

Springer-Lehrbuch

Oliver Deiser

Einführung in die Mengenlehre

Die Mengenlehre Georg Cantors und ihre
Axiomatisierung durch Ernst Zermelo

3., korrigierte Auflage

 Springer

Priv.-Doz. Dr. Oliver Deiser
Freie Universität Berlin
Fachbereich Mathematik
Arnimallee 6, Raum 211
14195 Berlin
Deutschland
deiser@mi.fu-berlin.de

ISSN 0937-7433

ISBN 978-3-642-01444-4

e-ISBN 978-3-642-01445-1

DOI 10.1007/978-3-642-01445-1

Springer Heidelberg Dordrecht London New York

Die Deutsche Nationalbibliothek verzeichnet diese Publikation in der Deutschen Nationalbibliografie; detaillierte bibliografische Daten sind im Internet über <http://dnb.d-nb.de> abrufbar.

Mathematics Subject Classification (2000): 03-01, 03-03, 01-01

© Springer-Verlag Berlin Heidelberg 2002, 2004, 2010

Dieses Werk ist urheberrechtlich geschützt. Die dadurch begründeten Rechte, insbesondere die der Übersetzung, des Nachdrucks, des Vortrags, der Entnahme von Abbildungen und Tabellen, der Funksendung, der Mikroverfilmung oder der Vervielfältigung auf anderen Wegen und der Speicherung in Datenverarbeitungsanlagen, bleiben, auch bei nur auszugsweiser Verwertung, vorbehalten. Eine Vervielfältigung dieses Werkes oder von Teilen dieses Werkes ist auch im Einzelfall nur in den Grenzen der gesetzlichen Bestimmungen des Urheberrechtsgesetzes der Bundesrepublik Deutschland vom 9. September 1965 in der jeweils geltenden Fassung zulässig. Sie ist grundsätzlich vergütungspflichtig. Zuwiderhandlungen unterliegen den Strafbestimmungen des Urheberrechtsgesetzes.

Die Wiedergabe von Gebrauchsnamen, Handelsnamen, Warenbezeichnungen usw. in diesem Werk berechtigt auch ohne besondere Kennzeichnung nicht zu der Annahme, dass solche Namen im Sinne der Warenzeichen- und Markenschutz-Gesetzgebung als frei zu betrachten wären und daher von jedermann benutzt werden dürften.

Einbandentwurf: WMXDesign GmbH, Heidelberg

Printed on acid-free paper

Springer ist Teil der Fachverlagsgruppe Springer Science+Business Media (www.springer.com)

für Caroline, Thalia und Larina



Die mathematische Theorie des Unendlichen

*Meine Gedankenwelt, d. h. die Gesamtheit S aller
Dinge, welche Gegenstand meines Denkens sein
können, ist unendlich.*

(Richard Dedekind in „Was sind und was sollen die Zahlen?“)

Inhalt

Vorwort	5
Historischer Überblick	10
1. Abschnitt Einführung	13
1. Mengen	15
2. Zwischenbetrachtung	43
3. Abbildungen zwischen Mengen	48
4. Größenvergleiche	64
5. Der Vergleichbarkeitssatz	81
6. Unendliche Mengen	91
7. Abzählbare Mengen	109
8. Überabzählbare Mengen	120
9. Mengen der Mächtigkeit der reellen Zahlen	130
10. Die Mächtigkeit der Potenzmenge	145
11. Die Kontinuumshypothese	149
12. Kardinalzahlen und ihre Arithmetik	160
13. Paradoxien der naiven Mengenlehre	183
Biographie von Georg Cantor	195
2. Abschnitt Ordnungen und Mengen reeller Zahlen	201
1. Transfinite Operationen	203
2. Lineare Punktmengen	207
3. Wohlordnungen	222
4. Der Fundamentalsatz über Wohlordnungen	230
5. Der Wohlordnungssatz	238
6. Ordinalzahlen	250
7. Transfinite Induktion und Rekursion	269
8. Typen linearer Ordnungen und ihre Arithmetik	284
9. Große Teilmengen und große Kardinalzahlen	306
10. Die Ordnungstypen von $\mathbb{Q}$ und $\mathbb{R}$	341
11. Der Satz von Cantor-Bendixson	360
12. Die Mächtigkeiten abgeschlossener Mengen	376
13. Die Vielheit aller Ordinalzahlen	401
Biographie von Felix Hausdorff	407

3. Abschnitt Die Basisaxiome der Mengenlehre	415
1. Das Axiomensystem ZFC	417
2. Die Sprache der Mengenlehre	444
3. Mengen und Klassen	468
Biographie von Ernst Zermelo	479
4. Abschnitt Anhänge	483
1. Liste der ZFC-Axiome	485
2. Lebensdaten der „dramatis personae“	486
3. Die wichtigsten Arbeiten von Cantor, Hausdorff und Zermelo	487
4. Zeittafel zur frühen Mengenlehre	499
5. Literatur	508
6. Notationen	542
7. Personenverzeichnis	545
8. Sachverzeichnis	547

Vorwort

Die Mengenlehre fängt bei nichts an. Als Basiswissen genügt die Intuition über Menge und Element, die fast jeder schon mitbringt und die gegebenenfalls leicht erweckt werden kann. Mit Hilfe weniger elementarer Konzepte läßt sich eine reiche mathematische Theorie begründen, und es lassen sich darin schnell tiefgreifende und zum Teil im Rahmen der Theorie unlösbare Fragen aufstellen.

Dieser Text will eine Einführung in die faszinierende Welt der unendlichen Mengen geben. Er ist gedacht für Studenten der Mathematik, Informatik und Philosophie in den ersten Semestern, insbesondere für solche, die mit der rudimentären Behandlung der Mengenlehre in den Anfängervorlesungen unzufrieden sind. Darüber hinaus ist das Buch geschrieben für jeden Mathematiker, der sich gerne die Sache mit der Mengenlehre noch einmal von Grund auf erklären lassen möchte oder eine Wissenslücke zu schließen sucht. Und auch der interessierte Laie und der Schüler nach oder während der Abiturzeit kann, so ist zu hoffen, dieses Buch mit Gewinn lesen.

Die Mengenlehre ist die Untersuchung von Ordnung und Größe in der Mathematik, ihre Wurzeln sind die Theorien der Wohlordnungen und der Mächtigkeiten. Letztere vorzuziehen eignet sich für einen einführenden Abschnitt nicht nur aus historischen Gründen, sondern auch deshalb, weil sie sich unmittelbar aus dem Funktionsbegriff entwickeln läßt. Allerdings sind zwei wichtige Sätze alles andere als einfach zu beweisen: Der Satz von Cantor-Bernstein und der Vergleichbarkeitssatz von Zermelo. Wir geben vollständige Beweise dieser beiden Sätze, sodaß der Anfänger auch härtere Brocken vorfinden wird.

Um die zentralen Objekte und Ideen vor dem geistigen Auge sehen zu können, ist die axiomatische Behandlung zunächst nicht notwendig, und hätte, gleich zu Beginn präsentiert, einen unangenehmen ad-hoc-Charakter. Erst die Paradoxien und metamathematische Fragen machen eine axiomatische Begründung notwendig.

Im ersten Abschnitt werden also die ersten Schritte der Theorie der Mächtigkeit rein aus der naiven Intuition des Mengenbegriffs heraus entwickelt und kommentiert. Vieles, was andernorts vielleicht verlorenging und zu schnell abgehandelt wird – die Ergebnisse aus den Geburtsjahren einer Theorie lassen sich ein Jahrhundert später immer glatt, sauber und schnell präsentieren – lebt hier noch einmal auf. Der Leser soll ein Gefühl für die Begriffe bekommen, ein solides Verständnis von dem, „worum es geht“. Wir setzen im ersten Abschnitt schamlos die natürlichen und die reellen Zahlen als gegeben und bekannt voraus.

Die übrigen Definitionen aber greifen nur auf bereits Definiertes zurück. Der Funktionsbegriff z. B. ist nicht mehr naiv, sondern streng mengentheoretisch, das heißt er wird auf die Relation *a ist Element von b* zurückgeführt. Die erzielten Kenntnisse lassen sich in einen axiomatischen Aufbau dann leicht und ohne große Wiederholungen einbinden.

Am Ende des ersten Abschnitts besprechen wir die in der naiven Mengenlehre auftretenden Paradoxien. Ein genauerer Blick auf die Fundamente wird also notwendig – und die Ergebnisse der naiven Mengenlehre zeigen, daß es der Mühe wert ist.

Bevor wir uns aber der axiomatischen Entwicklung zuwenden, behandeln wir in einem zweiten Abschnitt die beiden anderen großen Themenfelder der Cantorschen Forschung: Ordnungen und Teilmengen reeller Zahlen. Diese beiden auf den ersten Blick verschiedenen Gebiete sind historisch so stark miteinander verwoben, daß eine gemeinsame Behandlung am natürlichsten erschien; die Ordinalzahlen entdeckte Cantor bei der Untersuchung von Häufungspunkten von Teilmengen des Kontinuums. Nach einer Einführung in die Theorie der Ordnungen und insbesondere der Wohlordnungen untersuchen wir den transfiniten Prozeß der Ableitung einer Punktmenge sowie Größe und Struktur perfekter Mengen. Auf der Seite der Ordnungstheorie analysieren wir über die Wohlordnungen hinaus auch die Ordnungstypen der rationalen und der reellen Zahlen, und erhalten einen rein ordnungstheoretischen Beweis für die Überabzählbarkeit des Kontinuums. Die Antinomie der „Menge aller Ordinalzahlen“, die wir am Ende des zweiten Abschnitts diskutieren, führt schließlich wieder die Notwendigkeit einer genaueren Analyse des Mengenbegriffs vor Augen.

Im dritten Abschnitt stellen wir dann die heute am häufigsten verwendete Axiomatik der Mengenlehre vor, die Zermelo-Fraenkel-Axiomatik ZFC. Zunächst formulieren wir die Axiome in der üblichen mathematischen Umgangssprache und ziehen elementare Folgerungen. Anschließend besprechen wir die formale Sprache der Mengenlehre, formulieren die Axiome von ZFC in dieser Sprache und diskutieren Mengen und Klassen.

Im Anhang werden die wichtigsten Arbeiten von Georg Cantor, Felix Hausdorff und Ernst Zermelo kurz referiert, weiter findet der Leser die Lebensdaten der im Text häufiger genannten Mathematiker, sowie eine Tafel der ZFC-Axiome. Ein umfangreiches Literaturverzeichnis von Originalarbeiten, Lehrbüchern und historisch-philosophischen Texten soll die weitere Erkundung der Mengenlehre in verschiedene Himmelsrichtungen erleichtern.

Gelegentlich wird auf einen zweiten Teil des Buches verwiesen. In diesem Band wird die axiomatische Mengenlehre dann weiter systematisch entwickelt, und einige faszinierende moderne Erkenntnisse und Entwicklungen werden im Überblick vorgestellt. Bis zum Erscheinen des zweiten Bandes müssen wir den Leser auf die Literatur verweisen.

Eingestreut in den Text sind Übungsaufgaben verschiedenen Schwierigkeitsgrades. Ihre Kenntnisnahme ist wichtig, ihre Lösung nützlich für das Verständnis des übrigen Textes. Zu manchen Übungsaufgaben findet der Leser Hinweise in eckigen Klammern.

An diese Einleitung schließt sich ein historischer Überblick an, am Ende der ersten drei Abschnitte findet der Leser Skizzen der Biographien von Cantor, Hausdorff und Zermelo, den Hauptpersonen des Buches. Weiter sind Auszüge aus Originalarbeiten verschiedenster Art in den Text eingewoben, die dem Stoff eine weitere Dimension geben und ihn historisch befestigen. Der Schwerpunkt liegt aber auf den mathematischen Inhalten, die aus der heutigen Perspektive betrachtet und gewichtet werden, und die Darstellung folgt nicht in allen Punkten der geschichtlichen Entwicklung. Ein zuweilen eingestreuter Originalton ist aber für den Leser vielleicht eine willkommene Abwechslung und hoffentlich Motivation auch zur weitergehenden Lektüre.

Vorab sind vielleicht auch noch einige satztechnische und stilistische Bemerkungen angebracht.

Definitionen, Sätze, Korollare, Beweise, Übungen sowie *Originalzitate* innerhalb der Kapitel verlaufen auf Einzug und sind oftmals in variabler Satzbreite notiert. Sie sollen dadurch zu Einheiten werden, die sich vom Fließtext im Blocksatz abheben. Das Ende eines Beweises — ist zudem mit einem kleinen linken Randstrich gekennzeichnet.

In Beweisen wird eine *Annahme*, die widerlegt werden soll, immer kursiv gesetzt, wie auch der schließlich erhaltene *Widerspruch*. Innerhalb der Beweise werden oft neue Zeilen begonnen und wichtige Objekte werden an einer für das Auge gut auffindbaren Stelle definiert, indem sie z. B. durch jeweils einen halben Zeilenvorschub flankiert werden. Andererseits sind die Beweise kurz gehalten, und die dort verwendete Sprache ist elliptisch und karg. Das Ziel bei alledem ist, ein Lesen zu unterstützen, das Argumente nicht Schritt für Schritt abhaken will, sondern die innere Anschauung und das Wissen um die verwendeten Objekte zu vermehren trachtet, und dabei gleichzeitig die Struktur der Argumentation im Auge behält. Hierbei wird vom Leser ein gewisses Maß an Eigenarbeit und Ergänzung erwartet. Ein Beweis soll immer mehr vermitteln als ein Gefühl der Korrektheit.

Auf eine Numerierung der Sätze wurde bewußt verzichtet. Dafür sind aber viele Sätze mit Namen versehen, unter denen sie dann im weiteren Verlauf herangezogen werden. Ein Ausdruck „nach dem Satz oben“ bezieht sich in der Regel auf ein unmittelbar zuvor bewiesenes Resultat.

Zitate aus der Originalliteratur innerhalb eines Kapitels werden durch zwei waagrechte Linien eingeschlossen, solche am Ende des Kapitels haben lediglich eine Anfangslinie und insgesamt eine etwas andere Form. Diese Schlußzitate sollen ein Kapitel abrunden und ergänzen, während die Zwischenzitate einen Gegenstand betreffen, der in ihrer unmittelbaren Umgebung behandelt wird. In Zitaten kennzeichnen wie üblich eckige Klammern [...] Ergänzungen des Zitierenden. Das Schriftbild des Originals wurde nach Möglichkeit getreu wiedergegeben, insbesondere gilt das für *Kursivstellungen* und Sperrungen. Da die Zitate in Anführungszeichen stehen, werden die Anführungen innerhalb der Zitate zu einfachen Anführungen, also „...“ zu ‚...‘.

Der Hauptteil des Buches ist in drei Abschnitte eingeteilt, und diese jeweils in Kapitel. Innerhalb der Kapitel gibt es dann drei Typen von Sektionen, die durch

doppelt, einfach oder nicht unterstrichene Überschriften gekennzeichnet sind, und dadurch hierarchisch angeordnet werden. Der bei weitem häufigste Fall ist der zweite, also ein Unterkapitel, das durch eine einfach unterstrichene Überschrift eingeleitet wird.

Es wurde versucht, einigen klassischen Ansprüchen der Typographie gerecht zu werden: Vermeidung von „Hurenkindern“, also Zeilenresten am Seitenbeginn; korrekte Zwischenräume, etwa z. B. statt z.B. oder *wird?* statt *wird?*; richtige Anführungszeichen und Gedankenstriche; Vermeidung von mehr als drei aufeinanderfolgenden Trennungen, u. s. w. Dem Autor kam dabei das hochentwickelte Satzsystem *TypoScript* entgegen, mit dessen Hilfe er dieses Buch verwirklichen konnte. Dem Leser, der sonst immer nur L^AT_EX-Produkte zu sehen bekommt, ist dies hoffentlich eine willkommene Abwechslung. (Dank gilt an dieser Stelle meinen Eltern und darüber hinaus meinem Bruder Klaus, der nicht nur die wesentlichen Teile von *TypoScript* programmiert hat, sondern sich auch nicht vorhandene Zeit nahm, meine Makros zu reparieren.)

Entgegen dem traditionellen und auch heute üblichen Vorgehen, alle mathematischen Symbole kursiv zu setzen, behandeln wir diese – auch in Zitaten – gleichberechtigt, und schreiben etwa: „Sei f eine Funktion ...“ anstelle von „Sei f eine Funktion ...“. Kursivstellungen sind Auszeichnungen innerhalb eines Textes, und ob man einen Text zu einem Großteil aus Auszeichnungen bestehen lassen will, scheint zumindest eine Geschmacksfrage zu sein. (Ältere mathematische Schriften enthalten bei weitem mehr Text und weniger mathematische Symbole als zeitgenössische.)

Das Buch folgt der „alten“ Rechtschreibung des Deutschen.

Die verwendete Hauptschrift ist die Janson, zusammen mit der *Janson kursiv* und **Janson halbfett**. Sie ist ebenso klassisch und modern wie der Inhalt, den dieses Buch zu übermitteln hofft.

Besonderheiten der zweiten Auflage

Für die zweite Auflage wurde das Buch kapitelweise überarbeitet und dabei an vielen Stellen korrigiert und erweitert (von etwa 330 auf nun 550 Seiten). Größere Veränderungen gegenüber der ersten Auflage sind: die frühe Isolation eines Maximalprinzips (1.5), eine ausführlichere Behandlung verschiedener Unendlichkeitsdefinitionen (1.6), ein Kapitel über Kardinalzahlarithmetik im ersten Abschnitt, einschl. eines Beweises des Multiplikationssatzes und des Satzes von König-Zermelo (1.12), eine ausführlichere Diskussion der Paradoxien (1.13), eine Definition der Borel-Hausdorff-Hierarchie (2.7), verschiedene Beweise des Multiplikationssatzes mit Hilfe der Wohlordnungstheorie (2.8), ein einführendes Kapitel über Konfinalitäten, Filter und die frühen großen Kardinalzahlen „unerreichbar“, „Mahlo“ und „meßbar“ (2.9), Diskussion der Hausdorffschen Residuen-Operation (2.11), ausführlichere Diskussion von Zermelos Zahlreihe Z_0 (3.1), vollständige Angabe eines logischen Kalküls für die Sprache der Mengenlehre auf der Metaebene (3.2), kurze Diskussion von Axiomensystemen mit echten Klassen (3.3). Insgesamt sollte der Orientierungszeitraum 1870 – 1930 besser abgedeckt werden.

Die Zitate aus den Originalarbeiten wurden überprüft und an vielen Stellen korrigiert. Die historischen Anmerkungen wurden ergänzt und verbessert. Sie bleiben Anregung und bescheidener Einblick in die Zeiten der Vergangenheit, bekanntlich ein Buch mit sieben Siegeln. Sie sind Elemente der freien Erzählung eines Mathematikers und sind sicher nicht bis ins Letzte urkundlich und mathematikgeschichtlich belastbar.

Den Biographien und den Werkzusammenfassungen von Georg Cantor und Ernst Zermelo sind nun solche von Felix Hausdorff zur Seite gestellt. Ich hoffe, daß das Buch dadurch zusätzliche Symmetrie und innere Stabilität erhält. Das Triumvirat, das das frühe ϵ -Imperium regiert, deckt Geist und Form der Mengenlehre perfekt ab. Eine originellere und kontrastreichere Figurenkonstellation gibt es allenfalls noch im Theater.

Die knappe historische Einführung gleich zu Beginn des Buches wird jetzt durch eine ausführlichere Zeittafel zur Geschichte der frühen Mengenlehre ganz am Ende ergänzt. Das Literaturverzeichnis ist angewachsen und erscheint nun in kommentierter Form.

Auch andere hebräische Buchstaben als das Aleph sind nun in hebräischer Schrift wiedergegeben, was dem geheimnisvollen Charakter der Kardinalzahlarithmetik nur gerecht wird. Der Leser findet ein vollständiges hebräisches, griechisches, Fraktur- und Skriptur-Alphabet im Verzeichnis der Notationen.

Meinen Dank möchte ich den vielen Lesern aussprechen, die mir ihre Reaktionen mitgeteilt haben, oft verbunden mit Hinweisen auf Ungenauigkeiten und mit Vorschlägen zur Ergänzung. Die breite Herkunft der Leser war ein Grund mehr, die Linien des Buches unverändert beizubehalten.

Ich konnte all die mathematischen, stilistischen, didaktischen, historischen, philosophischen und typographischen Ideale, die mir vorschweben, wieder nur in Ansätzen realisieren. Ich glaube aber mit dieser zweiten Auflage dem Ziel, dem neugierig Zuhörenden eine vielseitige und spannende Geschichte zu erzählen, einen Schritt näher gekommen zu sein. Die Erzählung versucht verschiedene Aspekte des überlieferten Stoffs darzustellen, ohne dabei allzuoft in einen historisch - tragisch - komisch - pastoralen Stil zu verfallen. „Wer vieles bringt, wird manchem etwas bringen“ bleibt Motto bloß des Theaterdirektors, nicht des Autors. Wenn es am Ende zutrifft, so soll es aber recht sein.

München, im Januar 2004,

Oliver Deiser

Für die dritte Auflage wurde der Text erneut an vielen Stellen verbessert und korrigiert. Mein Dank gilt hier allen Lesern und ihren Zuschriften, vor allem aber Herrn Heinz Jaskolla.

Der Fortsetzungsband „Axiomatische Mengenlehre – Die Architektur von ZFC und die Unabhängigkeit der Kontinuumshypothese“ wird voraussichtlich im Frühjahr 2011 erscheinen.

Berlin, im Juli 2009,

Oliver Deiser

Historischer Überblick

Die Mengenlehre wurde im letzten Viertel des 19. Jahrhunderts von Georg Cantor entwickelt. Entgegen vorherrschenden Dogmen über den Umgang mit unendlichen, „fertigen“ Gesamtheiten, schuf er in einem gewaltigen Kraftakt die transfiniten Zahlen und das Konzept der Mächtigkeit oder Größe einer unendlichen Menge. Er entdeckte die Überabzählbarkeit der reellen Zahlen, das Kontinuumsproblem und untersuchte hinsichtlich einer Lösung des Problems die reellen Zahlen unter völlig neuartigen Gesichtspunkten.

Allerdings zeigte sich, daß man vorsichtig im Umgang sehr großer Gesamtheiten sein mußte. Georg Cantor war mit diesem Phänomen vertraut, aber er äußerte sich hierzu in seinen Veröffentlichungen nur marginal. Ernst Zermelo, Cesare Burali-Forti und Bertrand Russell fanden um die Jahrhundertwende Widersprüche der uneingeschränkten Mengenbildung:

*„Zu jeder Eigenschaft existiert die Menge aller Objekte,
auf die diese Eigenschaft zutrifft.“*

Dieses sogenannte naive Komprehensionsprinzip ist nicht haltbar. Nun riefen aber nicht die mathematischen Ideen Cantors, die er mit sicherer innerer Anschauung entwickelt hatte, die Unstimmigkeiten hervor, sondern verantwortlich hierfür war allein der unreflektierte Rahmen, in welchem die Mengenlehre damals stattfand. David Hilbert rief eine genauere Untersuchung der Grundlagen der Mathematik ins Leben.

Ernst Zermelo löste 1908 das Problem axiomatisch durch die Angabe eines Systems, das sorgfältig die Existenz bestimmter Mengen und die Bildung von Mengen aus anderen Mengen beschreibt. In der Folge wurde diese Axiomatik von Ernst Zermelo noch um zwei Axiome ergänzt, das Ersetzungsschema von Abraham Fraenkel (1922) und das Fundierungsaxiom von John von Neumann (1925) und Ernst Zermelo (1930). Weiter wurde die verwendete Sprache präzisiert, in der die Axiome formuliert werden und die den Begriff der „Eigenschaft“ einer Menge festlegt (Thoralf Skolem 1922). Das entstehende System aus Sprache und Axiomen, die Zermelo-Fraenkel-Axiomatik ZFC, wird heute

zumeist als Rahmen für die Mengenlehre verwendet. Die Widersprüche verschwinden in diesem System in natürlicher Weise, und alle Ideen Cantors leben darin in ihrer ursprünglichen Schönheit fort.

Es zeigte sich, daß der neue Rahmen der axiomatischen Mengenlehre groß genug war, um alle Objekte der Mathematik – Zahlen aller Art, Funktionen, geometrische Gebilde usw. – darin interpretieren zu können, d. h. es existiert eine auf dem Mengenbegriff basierende Definition dieser Begriffe, die alle erwünschten und in der Mathematik benötigten Eigenschaften der Objekte bereitstellt. Die Mengenlehre eignet sich damit als Grundlagendisziplin für die Mathematik selbst, und sie ist in ihrer universellen Fähigkeit zur Interpretation mathematischer Konstrukte bislang konkurrenzlos.

Ungelöst blieb allerdings Cantors Kontinuumsproblem, das die Frage stellt, ob die reellen Zahlen in der Hierarchie der Mächtigkeiten unmittelbar nach den natürlichen Zahlen erscheinen. Cantor sah lange Zeit zuversichtlich einer Lösung entgegen, scheiterte aber an diesem Problem, das dann David Hilbert auf dem Mathematikerkongreß in Paris 1900 an die erste Stelle seiner berühmten Liste von 23 offenen Fragen für das kommende neue Jahrhundert setzte. Schließlich „lösten“ Kurt Gödel (1938) und Paul Cohen (1963) das Problem, indem sie zeigten, daß es unlösbar war. Die Beweismethoden von Gödel und Cohen waren allgemein genug, um eine Fülle von ähnlichen Resultaten folgen zu lassen. Das Phantom der Unbeweisbarkeit oder Unabhängigkeit von mathematischen Aussagen war, nachdem Gödel seine Existenz bereits in den dreißiger Jahren abstrakt bewiesen hatte, real und greifbar geworden.

Mit der Entdeckung der Unabhängigkeit der Kontinuumshypothese beginnt die zweite Phase der Geschichte der Mengenlehre, in der Erweiterungen der ZFC-Axiomatik um sogenannte große Kardinalzahlaxiome ein Zentrum des Interesses bilden – und in der zuletzt auch neue Einsichten in das Kontinuumsproblem gewonnen werden konnten, wenn auch eine Entscheidung über die Größe der reellen Zahlen immer noch nicht gefallen ist. Wir werden im zweiten Band einige Aspekte dieser zweiten Phase der Mengenlehre skizzieren.

Heute ist die Mengenlehre nicht nur Rahmen für die Mathematik, sondern selbst eine schillernde mathematische Theorie. Sie fasziniert nach wie vor durch die stille Schönheit ihrer ersten Konzepte und durch deren erstaunliche und anscheinend noch bei weitem nicht ausgelotete Reichweite und Tragfähigkeit. Ihre Verzweigungen sind vielfältig und subtil miteinander verwoben, ihre Geschichte ist bis in die allerjüngste Zeit voll von Überraschungen und reich an dramatischen Entwicklungen. „*The old lady*“ (*Sabaron Shelah*) hat, kurz gesagt, alles, was man von einer großen mathematischen Theorie verlangen darf.

1. Mengen

Menge und Element

Wir besitzen ein intuitives Verständnis des Begriffs „Menge“ und der Beziehung „a ist ein Element der Menge b“. Für „a ist ein Element der Menge b“ schreiben wir kurz „ $a \in b$ “.

Besonders in dieser Einführung stützen wir uns auf dieses naive Verständnis des Mengenbegriffs. Georg Cantor (1845 – 1918) hat in seiner letzten mengen-theoretischen Arbeit die folgende Zusammenfassung oder Beschreibung unserer Intuition formuliert:

„Unter einer ‚Menge‘ verstehen wir jede Zusammenfassung M von bestimmten wohlunterschiedenen Objekten m unserer Anschauung oder unseres Denkens (welche die ‚Elemente‘ von M genannt werden) zu einem Ganzen.“

(Georg Cantor, 1895a)

Dies ist keine mathematische Definition im heute üblichen Sinne – was genau ist eine „Zusammenfassung“ oder ein „Ganzes“? –, und dennoch beschreibt sie recht genau unsere Vorstellung von einer Menge. Und sie enthält eine bemerkenswerte Feinheit: Cantor betont den Akt der Zusammenfassung zu einem Ganzen, zu einem Objekt. Die Mengenbildung verläuft hiernach zweistufig: Zuerst wird eine Vielheit, eine Ansammlung, ein Bereich betrachtet, und in einem zweiten Schritt wird diese Vielheit zu einer Einheit zusammengefaßt. Cantor war lange vor seiner Definition völlig klar, daß man nicht alle Vielheiten zu einer Menge zusammenfassen kann, daß also der zweite objektbildende Schritt nicht in jedem Falle legitim ist. Wir kommen erst später auf diesen wichtigen Punkt zurück, denn der durch die Intuition gewiesene Weg läßt sich soweit verfolgen, bis die Grenzen des Mengenbegriffs sichtbar und erfahrbar werden.

Extensionalität und Iteration

Der Cantorschen Definition fügen wir noch ein *Gleichheitskriterium* hinzu. Man kann argumentieren, daß sich dieses Kriterium für die Gleichheit zweier Mengen aus Cantors Definition ableiten läßt.

Zwei Mengen sind genau dann gleich, wenn sie dieselben Elemente haben.

(Extensionalitätsprinzip)

Richard Dedekind (1831 – 1916), der einen Aufbau der Arithmetik mit Hilfe des Mengenbegriffs entwickelte, hat in seinem Buch „Was sind und was sollen die Zahlen?“ – ein Klassiker der mathematischen Abteilung der Weltliteratur – eine sehr ähnliche intuitive Beschreibung von „Menge“ gegeben und dabei das Extensionalitätsprinzip explizit notiert. Mengen heißen bei ihm Systeme.

Dedekind (1888): „Im Folgenden verstehe ich unter einem Ding jeden Gegenstand unseres Denkens... Es kommt sehr häufig vor, daß verschiedene Dinge $a, b, c, \dots$ aus irgend einer Veranlassung unter einem gemeinsamen Gesichtspunkte aufgefaßt, im Geiste zusammengestellt werden, und man sagt dann, daß sie ein System S bilden. Man nennt die Dinge $a, b, c, \dots$ die Elemente des Systems S , sie sind enthalten in S ; umgekehrt besteht S aus diesen Elementen. Ein solches System S (oder ein Inbegriff, eine Mannigfaltigkeit, eine Gesamtheit) ist als Gegenstand unseres Denkens selbst ein Ding; es ist vollständig bestimmt, wenn von jedem Ding bestimmt ist, ob es Element von S ist oder nicht *). Das System S ist daher dasselbe wie das System T , in Zeichen $S = T$, wenn jedes Element von S auch Element von T ist, und jedes Element von T auch Element von S ist...“

Die Fußnote *) bei Dedekind holen wir gleich nach!

Neben der Extensionalität hebt Dedekind hier einen weiteren fundamentalen Gesichtspunkt hervor: Die Mengenbildung liefert ein Ding, und damit können Mengen als Dinge die Elemente von anderen Mengen sein, und diese wiederum Elemente von wieder anderen Mengen, usw. Das Mengenkonzept ist seiner Natur nach *iterativ*, und die Mengenlehre erhält durch das sich aufschaukelnde Wechselspiel, daß jede Menge b , die rechts in „ $a \in b$ “ auftaucht, auch links in „ $b \in c$ “ auftauchen kann, sowohl Struktur als auch Flexibilität.

Der Dritte im Bunde sei Felix Hausdorff (1868 – 1942), für dessen Ausdruckstärke und Gedankenklarheit dieser Text häufig als Zeittunnel dienen wird. Er formuliert die Grundgedanken mehrere Jahrzehnte später so:

Hausdorff (1927): „Eine Menge entsteht durch Zusammenfassung von Einzeldingen zu einem Ganzen. Eine Menge ist eine Vielheit, als Einheit gedacht. Wenn diese oder ähnliche Sätze Definitionen sein wollten, so würde man mit Recht einwenden, daß sie idem per idem oder gar *obscurem per obscurius* definieren. Wir können sie aber als Demonstrationen gelten lassen, als Verweisungen auf einen primitiven, allen Menschen vertrauten Denkkakt, der einer Auflösung in noch ursprünglichere Akte vielleicht weder fähig noch bedürftig ist. Wir wollen uns mit dieser Auffassung begnügen und es als Grundtatsache hinnehmen, daß ein Ding M in eigentümlicher, nicht definierbarer Weise gewisse andere Dinge $a, b, c, \dots$ und diese wiederum jenes bestimmen; eine Beziehung, die wir mit den Worten ausdrücken: die Menge M besteht aus den Dingen $a, b, c, \dots$ “

Der kompakte Slogan „Vielheit, als Einheit gedacht“ verweist wieder auf die Möglichkeit der Iteration, und zudem auf die Zweistufigkeit des Vorgangs: Betrachtung einer Vielheit und Objektbildung.

Hausdorff betont wieder die Extensionalität des Begriffs: Zu einer Menge gehören Elemente, und die Elemente bestimmen „wiederum“ die Menge selbst. Es

gibt keine „roten“ oder „grünen“ Mengen, die genau die Zahlen 1, 2 und 3 als Elemente enthalten. Es gibt nur ein Ding, das aus 1, 2, 3 besteht.

Keine gute Vorstellung wäre es dagegen, eine Menge als „Summe ihrer Elemente“ zu betrachten. Die Menge b etwa, die nur die Menge a als Element hat, ist nach dem Extensionalitätsprinzip sicher nicht mit a identisch, wenn a selbst mehr als ein Element besitzt. Die „Summe der Elemente“ von b wäre aber a .

Auch heute gilt „Menge“ als ein nicht weiter definierter Grundbegriff – irgendwo muß man anfangen. An einer intuitiven Erläuterung kommt aber kein einführender Text vorbei, und zumeist ist es die Cantorsche Definition von 1895, die hierfür als Ausgangspunkt gewählt wird. Dies ist kein Zufall, und von Vorteil auch nicht nur aus rein historischen Gründen: In seiner Gesamtschau der Mengenlehre hatte Cantor neben einer herausragenden Intuition eine Unbefangenheit, die wir heute, formalistisch und axiomatisch geschult, kaum mehr erreichen können.

Selbstbestimmtheit und freie Begriffsbildung

Es gibt neben der Extensionalität des Mengenbegriffs und der Iterierbarkeit der Mengenbildung noch einen dritten ganz wesentlichen Aspekt, den man die Selbstbestimmtheit der Mengen nennen könnte. Hierzu liefern wir zuerst die den Satz „[Ein System] ist vollständig bestimmt, wenn von jedem Ding bestimmt ist, ob es Element von S ist oder nicht *)“ zierende Fußnote nach. Sie lautet:

Dedekind (1888): „*) Auf welche Weise diese Bestimmtheit zu Stande kommt, und ob wir einen Weg kennen, um hierüber zu entscheiden, ist für alles Folgende gänzlich gleichgültig; die zu entwickelnden allgemeinen Gesetze hängen davon gar nicht ab, sie gelten unter allen Umständen. Ich erwähne dies ausdrücklich, weil Herr Kronecker vor Kurzem (im Band 99 des Journals für Mathematik, S. 334 – 336) der freien Begriffsbildung in der Mathematik gewisse Beschränkungen hat auferlegen wollen, die ich nicht als berechtigt anerkenne; näher hierauf einzugehen erscheint aber erst dann geboten, wenn der ausgezeichnete Mathematiker seine Gründe für die Notwendigkeit oder auch nur die Zweckmäßigkeit dieser Beschränkungen veröffentlicht haben wird.“

Das ist nun inhaltlich wie historisch von großer Bedeutung. Leopold Kronecker (1823 – 1891) gehörte als angesehener Mathematiker zu den aktiven Gegnern der Cantorschen Mengenlehre und des Cantor-Dedekindschen Mengenbegriffs. Er war Mitbegründer des konstruktivistischen und intuitionistischen Zweiges der Mathematik, der sich von der klassischen, mengentheoretisch fundierten Mathematik dadurch unterscheidet, daß viele Dinge nicht erlaubt sind, etwa Existenzbeweise, die keine konkreten Beispiele oder Algorithmen mitliefern, oder der logische Schluß von *nicht nicht* A auf A für Aussagen A . Den Nachweis der *allgemeinen* „Notwendigkeit oder auch nur der Zweckmäßigkeit“ der Freiheitsberaubung ist dieser Zweig bis heute schuldig geblieben, und die klassische Mathematik kann mit ihrer scharfen Trennung der Begriffe *Existenz* und *Algorithmus* konstruktive Fragen innerhalb ihrer zollfreien Landschaften sehr gut behandeln, ohne ständig auf ein „Rasen betreten verboten“ zu stoßen.

Rückblickend erscheint heute die von konstruktiver Seite geförderte Feinanalyse des formalen mathematischen Beweisbegriffs am interessantesten zu sein, die zu charakterstarken, wenn auch etwas kauzigen Subsystemen der klassischen Logik geführt hat (intuitionistische Logik und Minimallogik, vgl. auch 3.2.). Vorausblickend findet das nicht gerade ideologiefreie konstruktive Erbe seine Katharsis vielleicht einmal in Anwendungen in der Informatik. Die Ergebnisse müssen letztendlich immer zeigen, was für bestimmte Dinge „zweckmäßig“ ist und was nicht. Unwahrscheinlich erscheint es heute, daß der konstruktive Rahmen den weitergefaßten klassischen Rahmen einmal ersetzen wird, wie es von vielen Apologeten des Nichtdürfens einmal vorgesehen war.

Inhaltlich besagt die Selbstbestimmtheit des Mengenbegriffs, daß wir eine Menge A bilden und mit ihr operieren können, ohne in allen konkreten Fällen Fragen der Form „ $Ist a \in A$?“ oder „*Hat A die und die Eigenschaft?*“ beantworten zu können. Es zeigt sich, daß in der allgemeinen mathematischen Praxis Bestimmtheitsorgen nicht auftauchen. Die Menge aller Primzahlen wird man als bestimmt ansehen, sobald man weiß, was eine Primzahl ist, die endlos strittige Menge aller sinnvollen Steuergesetze kommt dagegen in der Mathematik erst gar nicht vor. Innerhalb der formalen axiomatischen Mengenlehre werden dann letzte Zweifel an der Bestimmtheit von Mengenbildungen ausgeräumt, da diese kunstsprachlich genau geregelt werden. Man hätte in den sonnigen Breiten der üblichen Mathematik ein Streusalz gegen Glatteisbildung nicht nötig, aber für viele gewagtere Expeditionen der mathematischen Logik ist eine technische Zusatzausrüstung unerläßlich. Erst die Selbsterkenntnis führt zum Sündenfall und zum Verlust eines gleichförmigen Klimas.

Cantor schrieb bereits 1882 zu Fragen der Bestimmtheit und Wohldefiniertheit von Mengenbildungen:

Cantor (1882b): „Eine Mannigfaltigkeit (ein Inbegriff, eine Menge) von Elementen, die irgend welcher Begriffssphäre angehören, nenne ich *wohldefiniert*, wenn auf Grund ihrer Definition und in Folge des logischen Prinzips vom ausgeschlossenen Dritten [d. h. es gilt, entweder A oder nicht A für alle Aussagen A , scholastisch *tertium non datur*; ein Drittes gibt es nicht] es als *intern bestimmt* angesehen werden muß, *sowohl* ob irgend ein derselben Begriffssphäre angehöriges Objekt zu der gedachten Mannigfaltigkeit als Element gehört oder nicht, *wie auch* ob zwei zur Menge gehörige Objekte, trotz formaler Unterschiede in der Art des Gegebenseins einander gleich sind oder nicht.

Im allgemeinen werden die betreffenden Entscheidungen nicht mit den zu Gebote stehenden Methoden oder Fähigkeiten in Wirklichkeit sicher und genau ausführbar sein; darauf kommt es aber hier durchaus nicht an, sondern *allein* auf die *interne Determination*, welche in konkreten Fällen, wo es die Zwecke fordern, durch Vervollkommen der Hilfsmittel zu einer *aktuellen (externen) Determination* auszubilden ist.“

Wir haben also alle Freiheiten in der Mengenbildung, sofern nicht die Mengenbildungen selber widersprüchlich sind. Das ist zugegebenermaßen ein fast schon häretischer Gedanke – nur Gott kann alles, was sich nicht selbst widerspricht. (Den Zusatz der Widerspruchsfreiheit machte bereits die Scholastik, weshalb die bekannte Frage, ob der allmächtige G. auch einen Stein zu machen in der Lage ist, den er selber nicht aufheben kann, dort unproblematisch ist: Er

kann es nicht, was seiner Allmacht keinen Abbruch tut.) Mephisto wird uns nachrufen, ob uns bei unserer Gottähnlichkeit der freien Mengenbildung und der Erkenntnis dessen, was Menge ist und was nicht, nicht bange wird. Wir kümmern uns aber nicht um den Teufel.

Man kann die Selbstbestimmtheit der Mengen auch so beschreiben: Wir können die weite Welt erforschen, ohne alle Details unserer Umgebung zu kennen. Und die mathematische Walz auf dem Fuhrwerk der internen Determiniertheit erweist sich als fruchtbar: Naheliegende Fragen können wir oft dann lösen, wenn wir über einen viel weiter entfernten Raum etwas in Erfahrung gebracht haben. Die Mathematik wird durch die geschickte iterierte Bildung von Objekten, die zunächst viele nur intern determinierte Eigenschaften haben, nicht blockiert oder gefährdet, sondern angetrieben und gefördert. Die Kenntnis einiger Eigenschaften naheliegender Objekte reicht aus, um einige Eigenschaften entfernterer Objekte extern bestimmen zu können, aus denen dann neue Erkenntnisse über nahe Objekte gewonnen werden können. Das Abschneiden dieser Schleife würde einen Ergebnismangel nach sich ziehen, der nicht zu verschmerzen wäre.

Damit sind wir bei einer Vorstellung angelangt, zu der die freie Mengenbildung und die Selbstbestimmtheit ihrer Produkte in natürlicher Weise führt: Die Mengenbildung ist eher eine Mengenbenennung. Es gibt eine mathematische Welt außerhalb von uns, die es zu erforschen gilt, ganz so, wie die Entdecker unbekannte Länder erforscht haben, wie die Astronomie uns heute das Weltall Stück für Stück näher bringt, oder wie ein Philologe eine vergessene alte Sprache entziffert. Es gibt Fragen über Fragen, und dann gibt es plötzlich Antworten, die frei von Willkür gegeben werden, und weltbildende Wirkung haben können, ohne dabei als absolute Wahrheiten auftreten zu müssen. Sie bleiben einfach Entdeckungen. Diese Vorstellung einer Mengenwelt außerhalb unserer selbst, so naiv sie sein mag, gehört zu den fruchtbarsten Konstrukten. Als idealistische menschliche Vorstellung bleibt sie stets außerhalb der Mathematik, aber dieses „außerhalb von etwas“ ist gerade das, was sie nährt. Ein kühles nächtliches Weltall und ein romantischer Betrachter passen gut zusammen.

Wir kommen gleich noch einmal auf diese Idee einer platonischen Mengenwelt zurück. Zunächst wenden wir uns dem Wort „Menge“ und der recht komplizierten Entstehung seiner Bedeutung selber zu.

Die Genese von „Menge“

Der Term „Menge“ selbst wurde von Bernard Bolzano (1781 – 1848) geprägt, und Cantor hat ihn von Bolzano übernommen. Bolzano beschreibt Mengen umständlich über „Inbegriffe“. In den posthum „aus dem schriftlichen Nachlasse des Verfassers“ herausgegebenen „Paradoxien des Unendlichen“ heißt es:

Bolzano (1851): „Es gibt Inbegriffe, die, obgleich dieselben Teile [Elemente] A, B, C, D, ... enthaltend, doch nach dem Gesichtspunkte (Begriffe), unter dem wir sie so eben auffassen, sich als verschieden ... darstellen... Wir nennen dasjenige, worin der

Grund dieses Unterschiedes an solchen Inbegriffen besteht, die Art der Verbindung oder Anordnung ihrer Teile. Einen Inbegriff, den wir einem solchen Begriffe unterstellen, bei dem die Anordnung seiner Teile gleichgültig ist (an dem sich also nichts für uns Wesentliches ändert, wenn sich bloß diese ändert), nenne ich eine Menge...”

Bei Cantor sind nun erstaunlicherweise Mengen oft mit einer im Hintergrund stillschweigenden „inbegrifflichen Anordnung“ versehen, von der der intuitive Mengenbegriff heute, nicht zuletzt aufgrund Cantors eigener Umschreibungen, völlig frei ist. Ordnungen werden heute Mengen erst nachträglich aufgeprägt in Form von Relationen (s. u.), insbesondere also in Form von anderen ungeordneten Mengen. Cantor stellt sich das Gegebensein von Mengen oft geordnet vor, aber er arbeitet mit ihnen dann derart, daß nur ihre Extension in die Argumente eingeht. Wir kommen später auf diese, außerhalb von ordnungstheoretischen Überlegungen völlig stumme, Ordnung der Cantorsche Mengenvorstellung noch zurück (2.6). Sie ist interessant, aber für Cantors mathematische Ergebnisse irrelevant und bei der Lektüre seiner Werke nicht störend.

Neben „Menge“ wurde im 19. Jahrhundert auch mehr oder weniger gleichwertig verwendet: Mannigfaltigkeit, Gesamtheit, Inbegriff, Varietät, Klasse, Vielheit, System.

Dedekinds Wortwahl „System“ orientiert sich an der griechischen Tradition der Zahl als System von Einheiten oder Monaden (μονάδων σύστημα). Diese Vorstellung wird bereits Thales (~ 625 – 547 v. Chr.) zugeschrieben, explizit steht sie bei Nikomachos von Gerasa etwa 100 n. Chr. Euklid (~ 295 v. Chr.) spricht nicht von Systemen, dafür wird bei ihm der Akt der Zusammenfassung betont: Eine Zahl ist eine aus Einheiten zusammengesetzte Menge, τὸ ἐκ μονάδων συνκείμενον πλῆθος, wobei πλῆθος Menge, Anzahl, Vielheit bedeutet. Dessen Wurzel ist πολυ-, das im fremdsprachlich angereicherten Deutschen als Vorsilbe *poly-* überlebt hat. Die Monaden hat Gottfried Wilhelm Leibniz (1646 – 1716) zu neuen Ehren gebracht, und auch noch Cantor wird derartige Einheiten bei seiner Definition von Kardinalzahl und Ordinalzahl verwenden. Euklid und eine lange Tradition trösten den extensional „verdorbenen“ Betrachter dieser Definitionen dann nur wenig: Eine Menge, die aus ununterscheidbaren Einheiten zusammengesetzt ist, kann nicht mehr als ein Element haben.

In anderen Sprachen sind heute gebräuchlich: englisch *set*, französisch *ensemble*, italienisch *insieme*, spanisch *conjunto*, niederländisch das für deutsche Ohren hübsche *verzameling*, neugriechisch σύνολο. Im Mittelhochdeutschen sagte man *manec*, wenn man „viel, reichlich“ meinte, was die Zunge dann zu *manch* und *mannigfach* abgeschliffen hat (das englische *many* gehört ebenfalls hierher). Die *Menge* und die in der Mathematik heute sehr geometrisch aufgeladene *Mannigfaltigkeit* sind also etymologische Schwestern. Das Wort *Menge* selber stammt vom mittelhochdeutschen *menige*, gebildet aus *manec* wie etwa *Länge* aus *lang*. *Menge* hat dagegen mit *vermengen* etymologisch zum Glück nichts zu tun. Das „viel, reichlich“ hinter dem Begriff mag erklären helfen, warum lange Zeit Mengen, die gar kein Element oder nur ein einziges besitzen, als Sonderfälle betrachtet, isoliert oder sogar mit ihrem einzigen Element verwechselt wurden.

Der Leser möge diese etymologischen Ausflüge verzeihen, aber der Autor ist seit der Zeit, als er zum ersten Mal erfahren hat, daß der altgriechische Begriff für *Wahrheit* wörtlich genommen nichts anderes bedeutet als das *Unverborgene* vom Wert des Zurückhaltens überzeugt. Er wird sich bis zum Auftreten von *Kardinalzahlen* aber zurückhalten.

Das Sein und das Epsilon

Das Zeichen ϵ für die Elementbeziehung, später stilisiert zu $\in$, hat Giuseppe Peano (1858 – 1932) 1889 in einer lateinisch geschriebenen Arbeit eingeführt, in der sich auch die bekannten Dedekind-Peano-Axiome der Zahlentheorie finden:

Peano (1889): „IV. De classibus.

Signo K significatur classis, sive entium aggregatio.

*Signum ϵ significat est. Ita $a \epsilon b$ legitur *a est quoddam b*; $a \in K$ significat *a est quaedam classis*; $a \in P$ significat *a est quaedam propositio*.“*

Das kleine Epsilon geht hierbei auf $\epsilon\sigma\tau\acute{\iota}\nu$, altgriechisch für „er, sie, es ist“ zurück, $a \in b$ meint also „a ist ein b“, a ist eines derer von b, das „Sein des Seienden“ der Mengen sind also ihre Elemente.

Georg Cantor gebrauchte überraschenderweise keine Abkürzung für den Ausdruck „a ist Element von b“, und erst in den 20er Jahren des 20. Jahrhunderts setzte sich eine Abkürzung und die Schreibweise von Peano durch.

Hausdorff (1927): „Die fundamentale Beziehung eines Dinges a zu einer Menge A, der es angehört, bezeichnen wir mit G. Peano in Wort und Formel folgendermaßen:

a ist Element von A: $a \in A$.“

In seinen ein halbes Tausend Seiten starken „Grundzügen der Mengenlehre“ von 1914 kommt Hausdorff wie Cantor noch ohne eine Abkürzung für die Elementbeziehung aus.

Der Platonismus in der Mathematik

Eine frühere Umschreibung des Mengenbegriffs von Cantor lautete:

„Unter einer Mannigfaltigkeit oder Menge verstehe ich nämlich allgemein jedes Viele, welches sich als Eines denken läßt, d. h. jeden Inbegriff bestimmter Elemente, welcher durch ein Gesetz zu einem Ganzen verbunden werden kann, und ich glaube hiermit etwas zu definieren, was verwandt ist dem Platonischen εἶδος oder ἰδέα ...“

(Georg Cantor, 1883b, Anmerkung 1)

Hier ist der Zusatz „... welches sich als Eines denken läßt“ von großer Bedeutung. Dieser Hinweis wäre überflüssig, wenn sich *jede* Vielheit, *jeder* Bereich von

Objekten als Einheit denken läßt. (Wir betonen diese Feinheiten, weil fälschlicherweise oft der Cantorsche Mengenbegriff mit der inkonsistenten naiven Mengenlehre identifiziert wird.)

Die explizite Verbindung der Mengen mit der Platonischen Ideenlehre ist für die Vorstellung, was durch die Mengenlehre beschrieben wird, seit jeher sehr klar und fruchtbar: Mengen existieren als Ideen unabhängig von uns. Und dies gilt allgemein für die Objekte der Mathematik. Wir reden nicht über geometrische Figuren – Kreise, Dreiecke, Geraden –, die wir in den Sand zeichnen, sondern über diese Figuren an sich. Ebenso ist es mit den Zahlen, die nicht das sind, was wir auf einem Papier in bestimmter Art und Weise notieren, sondern die für sich in einer wohldefinierten Realität vorhanden sind. Und erst recht gilt dies für die Mengen, die wir erst bei der Beschäftigung mit mathematischen Objekten entdecken, und die sich dann als fundamental herausstellen. (Der Hinweis, daß die Mathematik abstrahiert, bringt nichts an Klarheit. Gerade als eine Operation muß die Abstraktion auf etwas verweisen, und was sollte das Zielobjekt anderes sein als ein Ding außerhalb der Sinnenwelt.)

Existieren die mathematischen Gebilde zwar in einer durch die bloße Sinneswahrnehmung nicht zugänglichen Welt, so haben sie doch oftmals ihre Abbilder in der erfahrbaren Realität, durch die wir sie entdecken können – bei Platon: durch die unsere Seele sich an sie erinnert. Die enge Beziehung der Mengenlehre mit der Metaphysik, der Erkenntnis dessen, was hinter der erfahrbaren Realität existiert, hat Cantor Zeit seines Lebens betont. Jeder Mathematiker, der sich fragt, was er eigentlich tut, kommt an diesen Fragen nicht vorbei. Und die platonisch aufgefaßte Mengenlehre bildet dasjenige Gedankengebäude, an dem sich alle anderen Antworten reiben und verglichen mit dem alle anderen bisher vortragenen umfassenden Konzepte doch bloß als Strohütte neben einem griechischen Tempel erscheinen – zumindest aus der Sicht des Platonikers. Zuweilen liest man, daß der Platonismus heute in der Mathematik und der Mengenlehre nicht mehr aktuell sei. Dies ist keineswegs der Fall. Nicht zuletzt hat er aus ästhetischen und didaktischen Gründen Priorität, und was kann man von einer Sicht der Dinge besseres sagen, als daß sie das Verständnis der Dinge fördert und begünstigt. Wir diskutieren den konträren formalistischen Standpunkt im dritten Abschnitt. Der Leser ist aufgerufen, sich sein eigenes Bild zu machen.

Unendliche Mengen

Unendliche Mengen bilden das Zentrum der Mengenlehre, und sie werden dort als „fertige Gesamtheiten“ angesehen.

Hausdorff (1914): „Eine Menge ist eine Zusammenfassung von Dingen zu einem Ganzen, d. h. zu einem neuen Ding. Man wird dies allerdings schwerlich als Definition, sondern nur als anschauliche Demonstration des Mengenbegriffs gelten lassen, die auf einfache Beispiele verweist: wie etwa die Menge der Einwohner einer Stadt, die Menge der Wasserstoffatome der Sonne. Diese beiden Mengen sind endlich, sie be-

stehen aus einer endlichen, die zweite freilich aus einer ungeheuer großen Anzahl von Gegenständen. Es ist das Verdienst Georg Cantors, auch unendliche, d.h. nicht endliche Mengen in den Kreis der Betrachtung gezogen und damit, über populäre Vorurteile und philosophische Machtansprüche hinwegschreitend, eine neue Wissenschaft, die Mengenlehre, begründet zu haben; denn eine bloße Theorie der endlichen Mengen wäre ja nichts weiter als Arithmetik und Kombinatorik. Die Menge der natürlichen Zahlen, die Menge der Punkte des Raumes sind die nächstliegenden Beispiele unendlicher Mengen ...“

Über die Existenzberechtigung „fertiger“ unendlicher Objekte gab es im 19. Jahrhundert eine zuweilen polemische und überhitzte Diskussion. Man unterscheidet hier zwei Konzepte: „aktual unendlich“ und „potentiell unendlich“. Das folgende Beispiel zweier Aussagen über die natürlichen Zahlen erläutert diese beiden Begriffe vielleicht am besten: „Die Menge der natürlichen Zahlen existiert als ein mathematisches Objekt“ (aktual unendlich) bzw. „Es gibt zwar keine größte natürliche Zahl, aber eine fertige Gesamtheit der natürlichen Zahlen existiert nicht“ (potentiell unendlich).

Cantor hat mehrmals betont, daß das potentiell Unendliche das aktual Unendliche (oder Transfinites) voraussetzt. Lesenswert ist hierzu die folgende Fußnote aus dem philosophischen Aufsatz „Mitteilungen zur Lehre vom Transfiniten“ aus dem Jahre 1887, eine Polemik gegen Johann Friedrich Herbart (1776 – 1841):

Georg Cantor (1887): „Nach Herbart ... soll der Begriff des Unendlichen ‚auf einer wandelbaren Grenze, welche in jedem Augenblick weiter fortgeschoben werden kann, bzw. muß‘ beruhen ... Ist es den Herren gänzlich aus der Erinnerung gekommen, daß, von den Reisen abgesehen, die in der Phantasie oder im Traume ausgeführt zu werden pflegen, daß, sage ich, zum sicheren Wandeln oder Wandern fester Grund und Boden sowie ein geebener Weg unbedingt erforderlich sind, ein Weg, der nirgends abbricht, sondern überall, wohin die Reise führt, gangbar sein und bleiben muß? ... Die weite Reise, welche Herbart seiner ‚wandelbaren Grenze‘ vorschreibt, ist eingestandenermaßen nicht auf einen endlichen Weg beschränkt; so muß denn ihr Weg ein unendlicher; und zwar, weil er seinerseits nichts Wandelndes, sondern überall fest ist, ein aktualunendlicher Weg sein. Es fordert also jedes potentiale Unendliche (die wandelnde Grenze) ein Transfinitum (den sicheren Weg zum Wandeln) und kann ohne letzteres nicht gedacht werden... Da wir uns aber durch unsere Arbeiten der breiten Heerstraße des Transfiniten versichert, sie wohlfundiert und sorgsam gepflastert haben, so öffnen wir sie dem Verkehr und stellen sie als eiserne Grundlage, nutzbar allen Freunden des potentialen Unendlichen, im besondern aber der wanderlustigen Herbartischen ‚Grenze‘ bereitwillig zur Verfügung; gern und ruhig überlassen wir die Rastlose der Eintönigkeit ihres durchaus nicht beneidenswerten Geschicks; wandle sie nur immer weiter, es wird ihr von nun an nie mehr der Boden unter den Füßen schwinden. Glück auf die Reise!“

Heute hat sich das aktual unendliche Konzept in der Mathematik erfolgreich durchgesetzt. Die Gefahr, daß der Begriff einer fertigen unendlichen Gesamtheit letztendlich widersprüchlich ist, besteht zwar weiterhin, und ein kritisches Bewußtsein ist sicher angebracht, zumal in der Natur die Endlichkeit – entge-

gen allem Anschein – vorzuherrschen scheint (vgl. auch die Diskussion und das Hilbert-Zitat im Kapitel 6 über unendliche Mengen). Andererseits wird nun seit fast hundert Jahren intensiv und erfolgreich in den meisten Teilgebieten der Mathematik mit unendlichen Objekten gearbeitet, und die Grundlagenforschung ist, im Gegensatz zur Situation im 19. und frühen 20. Jahrhundert, erwachsen geworden. Mit einem einfachen Widerspruch ist nicht zu rechnen, und das aus dem Konzept der aktual unendlichen Mengen heraus bewiesene Resultat „eine unendliche Menge existiert nicht“ oder „die Menge der reellen Zahlen existiert nicht“ wäre mit hoher Wahrscheinlichkeit keine Katastrophe, sondern vielmehr eine der tiefsten Erkenntnisse der Mathematik und auch des menschlichen Verstandes.

Gegen eine solche Götterdämmerung und die Widersprüchlichkeit der Theorie aktual unendlicher Mengen sprechen viele Argumente. Nicht zuletzt sind die Resultate der Mengenlehre von einer derart ergreifenden Schönheit, daß ein Zusammenbrechen des Gebäudes zusammen mit der Wahrung des Vertrauens, daß uns Natur und Verstand nicht an der Nase herum führen, schwer vorstellbar ist.

Die philosophische Frage nach der Unendlichkeit bleibt aber bestehen. Allerdings scheint es, daß sie auf einem gewissen Niveau nur auf dem Boden der mengentheoretischen Resultate diskutiert werden kann, und sich in Vagheiten verliert, wenn sie sich zu weit davon entfernt. Sicher der beste jüngere Beitrag zur Unendlichkeit ist das durch die mengentheoretische Forschung errichtete Gebäude der großen Kardinalzahlexiome, dessen Bedeutung für die Mathematik noch nicht abzusehen ist. Wir werden die ersten Stockwerke im zweiten Abschnitt erkunden.

Stellvertretend für die andere Seite sei Carl Friedrich Gauß (1777 – 1855) zitiert mit seinem Verdikt des aktual Unendlichen, das, so ernst es aus solchem Munde zu nehmen ist, doch zu jenen „Vorurteilen und Machtansprüchen“ gehört hat, die Cantor das Leben schwer gemacht haben:

„So protestiere ich zuvörderst gegen den Gebrauch einer unendlichen Größe als einer vollendeten, welches in der Mathematik niemals erlaubt ist. Das Unendliche ist nur eine façon de parler [Sprechweise], indem man eigentlich von Grenzen spricht, denen gewisse Verhältnisse so nahe kommen als man will, während anderen ohne Einschränkung zu wachsen gestattet ist.“

(Gauß an Heinrich Schumacher (1780 – 1850), 12. Juli 1831)

Diese Briefstelle kann man – allerdings nur mit viel Wohlwollen – auch als berechtigte Kritik daran lesen, mit den damals unzureichend definierten infinitesimalen Größen so zu rechnen, als wären sie wohldefinierte Objekte.

Zu „potentiell unendlich“ und „aktual unendlich“, und zur Genese der systematischen Untersuchung aktual unendlicher Mengen durch Georg Cantor schreibt Abraham Fraenkel (1891 – 1965) dann aus relativ großer zeitlicher Distanz:

Abraham Fraenkel (1959): „In der ‚klassischen‘ Mathematik des 19. Jahrhunderts tritt das Unendliche im allgemeinen in ‚potentieller‘ Form auf. Mittels dieses potentiellen Unendlichkeitsbegriffs haben A. L. Cauchy und seine Nachfolger in der ersten Hälfte des 19. Jahrhunderts die Infinitesimalrechnung streng begründet, dann K. Weierstraß,

G. Cantor und H. Méray in der zweiten Hälfte den Zahlbegriff, insbesondere die irrationalen Zahlen, und auf dieser Grundlage die Funktionentheorie aufgebaut. Es wird genügen ein ganz einfaches Beispiel zu geben: die Aussage $\lim_{n \rightarrow \infty} 1/n = 0$ (gelesen: der Limes (Grenzwert) von $1/n$, wenn n nach Unendlich strebt, ist Null (oder auch unendlichklein)) ist nichts als eine symbolische Verkürzung der Behauptung: der Quotient $1/n$, wo n eine natürliche Zahl bedeutet, kann mit *beliebiger* Genauigkeit an Null angenähert werden dadurch, daß n *genügend groß* gewählt wird. Offensichtlich ist in dieser Behauptung von Unendlichgroßem oder -kleinem nicht die Rede...

Ungeachtet gewisser tastender Ansätze einiger Mathematiker in den 70er und 80er Jahren, wie H. Hankel, A. Harnack, P. du Bois Reymond, ist es erst Georg CANTOR (1845 – 1918), der *Schöpfer der Mengenlehre*, gewesen, der das aktuelle Unendlich sorgfältig begründet, systematisch in Mathematik und Philosophie eingeführt und zur Grundlage einer eigenen Disziplin, eben der Mengenlehre, gemacht hat, die seit der Jahrhundertwende siegreich fast alle mathematischen Disziplinen infiltriert und weitgehend umgestaltet hat. Indes ist die Schöpfung der Lehre vom aktuellen Unendlich nicht etwa zielbewußt vom Anfang an beabsichtigt gewesen. Seit 1870 ging Cantor, der zeitlebens in Halle lehrte, von konkreten mathematischen Problemen der Theorie der reellen Funktionen aus, bei denen es auf die Unterscheidung endlich- oder unendlichvieler ‚Ausnahmepunkte‘ ankam, und rang sich nur allmählich, über eigene Hemmungen hinweg und dem heftigen Widerstand seiner mathematischen Zeitgenossen zuwider, zu einer allgemeinen revolutionären Begriffsbildung durch...

Mit dem Begriff „Mengenlehre“ ist oft zweierlei gemeint: Zum einen die Sprache der Mengenlehre samt ihrem aktual unendlichen Konzept und den elementaren Begriffen und Operationen, die sie der Mathematik als Grundwortschatz zur Verfügung stellt, und zum anderen die abstrakte mathematische Theorie, die die Struktur des weiten mengentheoretischen Universums zu ergründen sucht. Die Sprache der Mengenlehre lag Ende des 19. Jahrhunderts sicher in der Luft, und wurde vielerorts geschmiedet, etwa bei Richard Dedekind. Die Theorie ist Werk von Georg Cantor alleine, und trägt auch heute noch seine Handschrift. Darüber hinaus war er in der Formung der Sprache – als natürliche Folge der Herausbildung seiner Theorie – der begabteste Feinschmied. Die „Infiltrierung und Umgestaltung“ aller mathematischer Disziplinen, von der Fraenkel spricht, bezieht sich auf die Sprache der Mengenlehre und nicht auf die Mengenlehre als mathematische Theorie. Man kann auch nach Cantor noch sehr gut Mathematik betreiben, ohne viel von höheren Mächtigkeiten und ohne irgendetwas von Wohlordnungen zu wissen. Mit unendlichen Objekten haben die meisten Mathematiker dagegen tagtäglich zu tun. In diesem Buch bezieht sich der Ausdruck „Mengenlehre“ zumeist auf die mathematische Theorie, wobei die Theorie in offensichtlicher Weise ebenfalls von der Sprache der Mengenlehre Gebrauch macht.

Nach dieser Beschreibung der Intuition und ihrer formenden Ausgestaltung in einen extensionalen, iterativen und freien Mengenbegriff, zusammen mit einigen historischen und philosophischen Bemerkungen, die die Komplexität der zugehörigen Fragen vielleicht erahnen lassen, werden wir nun die ersten Erkundungen im Reich der mathematischen Mengen unternehmen. Zu Anfang einer mathematischen Theorie gibt es immer viele Definitionen. Aber wir müssen eine

Grundsprache und in ihr eine gewisse Geläufigkeit entwickeln, um über Mengen unangestrengt reden zu können. Und neue Worte helfen ja immer auch, sehen zu lernen, was es alles gibt.

Mengen aus mathematischen Objekten

Anstatt mit beliebigen „Objekten der Anschauung oder unseres Denkens“ wie in Cantors intuitiver Beschreibung von 1895 wollen wir uns hier nur mit Objekten der Mathematik beschäftigen. In „ $a \in b$ “ sollen also a und b mathematische Objekte bezeichnen.

Es wäre nicht nötig, neben den Mengen andere Objekte zuzulassen – sogenannte Grund- oder Urelemente –, denn es hat sich gezeigt, daß man alle in der Mathematik gebrauchten Objekte, insbesondere auch die natürlichen Zahlen, geeignet als Mengen interpretieren kann. Für die ersten zwei Abschnitte machen wir der Einfachheit halber die folgende Konvention. Eine gewisse Kenntnis der Grundzahlen setzen wir dabei voraus.

Konvention

Wir setzen für diesen und den nächsten Abschnitt die natürlichen, ganzen, rationalen und reellen Zahlen samt ihren üblichen Rechenoperationen und den natürlichen Ordnungsbeziehungen $<$ und $\leq$ als gegeben voraus.

Mathematische Objekte (innerhalb der ersten beiden Abschnitte)

- (i) Die mathematischen Grundobjekte (Urelemente) sind:
Natürliche Zahlen, ganze Zahlen, rationale Zahlen, reelle Zahlen.
- (ii) Jede Menge ist ein mathematisches Objekt.

Die mathematischen Objekte bestehen also aus den Grundobjekten und den Mengen. (Die Rechenoperationen und die Ordnungsbeziehung auf den Grundzahlen sind Funktionen und Relationen, die wir als Mengen von geordneten Paaren auffassen, siehe Kapitel 3.)

Wir halten noch fest: Ist b eine Menge und $a \in b$, so ist a ein Grundobjekt oder eine Menge.

Die bescheidene Auswahl der Grundobjekte in (i) geschah lediglich aus Gründen der Defintheit. Insbesondere die natürlichen und die reellen Zahlen werden für die ersten Schritte zur Erkundung unendlicher Mengen eine zentrale Rolle spielen. Wir diskutieren unten wichtige Eigenschaften dieser Zahlen. Der Leser, dem obige Konvention zu eng erscheint, kann zusätzlich beliebige Objekte der Mathematik betrachten, um sich Beispiele für Mengen und Mengen von Mengen zu konstruieren.

Im dritten Abschnitt werden wir auf Urelemente ganz verzichten. Daß alles aus dem Nichts entsteht ist zwar eine kulturgeschichtlich vertraute Idee, die Erzählung scheint dem Autor aber mit „Im Anfang waren die Zahlen ...“ viel flüssiger in Gang zu kommen als mit „Im Anfang war die leere Menge ...“.

Die Grundobjekte

Wir können die Grundobjekte zu Mengen zusammenfassen:

Definition

Wir setzen:

$\mathbb{N}$ = „die Menge der natürlichen Zahlen“,

$\mathbb{Z}$ = „die Menge der ganzen Zahlen“,

$\mathbb{Q}$ = „die Menge der rationalen Zahlen“,

$\mathbb{R}$ = „die Menge der reellen Zahlen“.

Da wir $\mathbb{N}$, $\mathbb{Z}$, $\mathbb{Q}$, $\mathbb{R}$ schlicht als gegeben voraussetzen, sind einige Bemerkungen angebracht.

Natürliche und ganze Zahlen

Die natürlichen Zahlen sollen die 0 als Element enthalten. Wir schreiben natürliche Zahlen wie üblich zumeist in Dezimalschreibweise:

0, 1, 2, 3, ..., 10, 11, ...

Eine ganze Zahl schreiben wir in der Form $+n$ oder $-n$ wobei n eine natürliche Zahl ist:

$+0, -0, +1, -1, +2, -2, \dots$

Es gilt $+0 = -0$.

Eine wichtige Eigenschaft der natürlichen Zahlen ist: Jede nichtleere Menge von natürlichen Zahlen hat ein kleinstes Element. Ist A eine Menge von natürlichen Zahlen, die mindestens ein Element enthält, so sei

$\min(A)$ = „das (eindeutig bestimmte) kleinste Element von A “.

[gelesen: Minimum von A].

Ist also z. B. A die Menge bestehend aus 8, 11, 5, 7, so ist $\min(A) = 5$.

Rationale Zahlen

Die rationalen Zahlen schreiben wir *entweder* als Brüche in der Form $+n/m$ oder $-n/m$, wobei n, m natürliche Zahlen sind mit $m \neq 0$ *oder* als endlichen oder periodisch endenden unendlichen Dezimalbruch in der Form

$\pm n, a_1 \dots a_k$ bzw.

$\pm n, b_1 \dots b_m a_1 \dots a_k a_1 \dots a_k a_1 \dots a_k \dots$,

wobei n, m, k natürliche Zahlen sind, und $0 \leq a_i < 10$ gilt für alle $1 \leq i \leq k$.

So gilt etwa:

$$-17/8 = -2,125,$$

$$1/3 = 0,333\dots,$$

$$1/7 = 0,142857142857142857\dots,$$

$$1/700 = 0,00142857142857142857\dots$$

Reelle Zahlen und kanonische Darstellung

Reelle Zahlen schreiben wir zumeist als Dezimalbruch $\pm n, a_1 a_2 a_3 \dots$, wobei n eine natürliche Zahl ist und $0 \leq a_i < 10$ gilt für alle $i \geq 1$.

Ist die Dezimaldarstellung einer reellen Zahl x von der Form

$$\pm n, a_1 \dots a_k 000 \dots,$$

so sagen wir, daß diese Dezimaldarstellung von x *trivial endet*. Jede reelle Zahl x – außer der Null! – hat eine *eindeutige nicht trivial endende Dezimaldarstellung*.

So gilt etwa:

$$\begin{aligned} 1,000 &= 0,999 \dots, \\ 1,1245000 \dots &= 1,1244999 \dots, \\ -1,01000 &= -1,00999 \dots \end{aligned}$$

Die nicht trivial endende Dezimaldarstellung von $x \in \mathbb{R}$, $x \neq 0$, bezeichnen wir als *die kanonische (unendliche) Dezimaldarstellung von x* . Weiter nennen wir $0,000 \dots$ die *kanonische (unendliche) Dezimaldarstellung der 0*.

Derartige Pedanterien sind in der Mengenlehre leider notwendig, da oftmals mit den Nachkommastellen von reellen Zahlen jongliert wird, und es hierfür auf eindeutige Darstellungen ankommt. Daß wir hier im Zweifel den unendlichen Darstellungen den Vorzug geben, geschieht einzig aus Gründen späterer Bequemlichkeiten.

Für jede natürliche Zahl $b \geq 2$ und jede reelle Zahl x existiert weiter eine b -adische Darstellung von x :

$$x = \pm n, a_1 a_2 a_3 \dots \text{ mit } 0 \leq a_i < b.$$

Dann ist $x = \pm (n + a_1/b + a_2/b^2 + a_3/b^3 + \dots)$.

Wie für die 10-adische, also die Dezimaldarstellung, existiert für jede reelle Zahl $x \neq 0$ eindeutig die nicht trivial endende *kanonische b -adische Darstellung von x* . Die 2-adische Darstellung heißt auch *Binärdarstellung*.

So ist z. B. $0,111 \dots$ die kanonische Binärdarstellung der 1.

Für alle $b \geq 2$ sei wieder $0,000 \dots$ die *kanonische b -adische Darstellung der 0*.

In diesem Buch brauchen wir neben der Dezimaldarstellung und der Binärdarstellung nur noch die 3-adische oder *Ternärdarstellung* (bei der Diskussion der Cantormenge).

Konvention

Wir identifizieren:

$$n \in \mathbb{N} \quad \text{mit} \quad +n \in \mathbb{Z},$$

$$+n \text{ bzw. } -n \in \mathbb{Z} \quad \text{mit} \quad +n/1 \text{ bzw. } -n/1 \in \mathbb{Q},$$

$$\pm n, a_1 \dots a_k \in \mathbb{Q} \quad \text{mit} \quad \pm n, a_1 \dots a_k 000 \dots \in \mathbb{R}.$$

Somit ist jede natürliche Zahl eine ganze Zahl, jede ganze Zahl eine rationale Zahl, und jede rationale Zahl eine reelle Zahl. (Periodische rationale Zahlen sind bereits vor dieser Konvention auch reelle Zahlen.)

Wir brauchen schließlich noch einige Notationen. Der *Betrag* $|x|$ einer reellen Zahl x ist definiert durch:

$$|x| = \begin{cases} x, & \text{falls } x \geq 0, \\ -x, & \text{falls } x < 0. \end{cases}$$

Die gleiche Notation verwenden wir später für Mengen, wo $|M|$ die Mächtigkeit einer Menge bezeichnet. Dies ist aber ungefährlich.

Sind $x_1, \dots, x_n$ reelle Zahlen, so bezeichnen wir mit $\min(x_1, \dots, x_n)$ die kleinste der Zahlen $x_1, \dots, x_n$. Analog bezeichnet $\max(x_1, \dots, x_n)$ die größte der Zahlen $x_1, \dots, x_n$. So ist z. B. $\min(0, -1) = -1$, $\max(2, 4, 3/2) = 4$.

Einfache Mengenbildungen

Wir bezeichnen Mengen mit lateinischen, griechischen, Fraktur-, Skriptur-, usw. Buchstaben (z. B. $a, b, N, M, \gamma, \Gamma, \alpha, \mathfrak{A}, \mathcal{A}, \mathcal{M}, \dots$). Welche Mengen diese Buchstaben bezeichnen, ist abhängig vom Kontext. Für bestimmte Mengen, die häufig auftauchen, gibt es feste, kontextunabhängige Zeichen, wie etwa das schlanke $\mathbb{N}$ für die Menge der natürlichen Zahlen.

Der Leser findet verschiedene Alphabete mit den Namen der Buchstaben im Notationsteil des Buches.

Viele von diesen variablen Bezeichnungen suggerieren einen bestimmten Bereich ihrer Bedeutung: Die Variablen n, m, k werden zumeist für natürliche Zahlen verwendet; ist von reellen Zahlen die Rede, so sind die Zeichen x, y, z erste Wahl; weiter ist A ein strukturell komplizierteres Objekt als a , und $\mathfrak{A}$ oder $\mathcal{A}$ ist noch komplizierter als A . (Warnung: In der Mengenlehre bedeuten häufig auch kleine Buchstaben reichhaltige Mengen.)

Das ständige Wiederholen gleicher Zeichen in ähnlichen Kontexten hat eine erstaunliche – und erwünschte! – psychologische Wirkung: Man vergleiche: „seien $n \in \mathbb{N}$ und $x_1, \dots, x_n \in \mathbb{R}$ “ mit dem formal gleichwertigen „seien $x \in \mathbb{N}$ und $n_1, \dots, n_x \in \mathbb{R}$ “. Irgendwann sind aber im Kopf alle Zeichen belegt, und somit sind Überschneidungen nicht zu vermeiden.

Den Ausdruck „ $a \in b$ “ lesen wir als:

„ a ist Element von b “, kurz „ a Element b “, oder auch „ a in b “.

Für „nicht $a \in b$ “ oder scholastischer „non($a \in b$)“ schreiben wir auch „ $a \notin b$ “. Wir vereinbaren zudem, daß $a \in b$ für alle a immer falsch und $a \notin b$ für alle a immer richtig ist, falls b keine Menge ist.

Wir können jede konkrete Liste mathematischer Objekte $a_1, \dots, a_n$ zu einer Menge zusammenfassen. Zur Bezeichnung verwenden wir die geschweiften Mengenklammern „ $\{$ “ und „ $\}$ “:

Definition (*direkte Angabe der Elemente; Einermengen, Paarmengen*)

Seien $n \in \mathbb{N}$, $n \geq 1$ und $a_1, \dots, a_n$ Objekte. Wir setzen

$\{a_1, \dots, a_n\} =$ „die Menge, die genau $a_1, \dots, a_n$ als Elemente enthält“.

Speziell heißen für Objekte a, b die Menge $\{a\}$ die *Einermenge von a* ,
und $\{a, b\}$ die (*ungeordnete*) *Paarmenge von a, b* .

Die Verwendung von geschweiften Klammern für die Notation von Mengen geht auf Georg Cantor und das Jahr 1895 zurück; er schrieb allerdings Mengen M in der heute irritierenden Form $M = \{m\}$, also M als Zusammenfassung von (vielen) Objekten m . Zuvor (ab 1874) verwendete Cantor auch Notationen der Form (m) , etwa (v) für die natürlichen Zahlen. Die Schreibweise der Definition oben hat dann Ernst Zermelo (1871 – 1953) eingeführt:

Zermelo (1908): „Die Menge, welche nur die Elemente $a, b, c, \dots, r$ enthält, wird zur Abkürzung vielfach mit $\{a, b, c, \dots, r\}$ bezeichnet werden.“

Nach Definition gilt für alle x :

$x \in \{a_1, \dots, a_n\}$ gdw $x = a_i$ für ein $i \in \mathbb{N}$ mit $1 \leq i \leq n$.

Die Abkürzung „gdw“ steht für „genau dann, wenn“, und meint dasselbe wie das schwerfällige „dann und nur dann“. Ein Ausdruck der Form „ A gdw B “ ist logisch gleichwertig mit „aus A folgt B , und aus B folgt A “.

Für die Bildung der Paarmenge ist $a \neq b$ nicht vorausgesetzt! Nach dem Extensionalitätsprinzip gilt:

$\{a\} = \{a, a\} = \{a, a, \dots, a\}$, $\{a, b\} = \{b, a\}$,

$\{a, b\} = \{a\}$ gdw $a = b$.

Allgemein können wir den folgenden nicht besonders spektakulären Sachverhalt konstatieren:

Übung

Seien $n, m \in \mathbb{N}$ und $a_1, \dots, a_n, b_1, \dots, b_m$ Objekte mit den Eigenschaften:

(α) Für alle $1 \leq i \leq n$ existiert ein $1 \leq j \leq m$ mit $a_i = b_j$.

(β) Für alle $1 \leq j \leq m$ existiert ein $1 \leq i \leq n$ mit $b_j = a_i$.

Dann gilt $\{a_1, \dots, a_n\} = \{b_1, \dots, b_m\}$.

Es genügt, wenn der Leser verinnerlicht, daß eine Menge $M = \{a, b\}$ nicht notwendig zwei Elemente haben muß. Es kann $a = b$ gelten, und dann ist M einelementig.

Weiter definieren wir:

Definition (*leere Menge*)

Wir setzen

$\emptyset =$ „die Menge, die kein Element enthält“.

$\emptyset$ heißt *die leere Menge* oder *die Nullmenge*.

Wir verwenden auch „ $\{ \}$ “ als Bezeichnung für die leere Menge, in Erweiterung der Schreibweise $\{ a_1, \dots, a_n \}$.

Cantor, Hausdorff und Zermelo haben noch das Symbol 0 für die leere Menge verwendet. André Weil (1906 – 1998) hat dann das $\emptyset$ -Zeichen aus den nordischen Sprachen eingeführt, und die Mathematiker mit ihrer Vorliebe für alle möglichen Alphabete haben dieses Zeichen dann in ihren festen Vorrat übernommen.

Die leere Menge kann Element einer anderen Menge sein. $M = \{ \emptyset \} = \{ \{ \} \}$ hat ein Element, nämlich die leere Menge. $M = \{ \emptyset, \{ \emptyset \} \}$ hat zwei Elemente: $\emptyset$ und $\{ \emptyset \}$ sind verschieden, wie man mit dem Extensionalitätsprinzip sofort sieht.

Hausdorff (1914): „Außer den Mengen, die Elemente haben, lassen wir auch eine Menge 0, die Nullmenge, zu, die kein Element hat; die Gleichung $A = 0$ bedeutet also, daß auch die Menge A kein Element hat, verschwindet, leer ist. Auch hierzu ist die analoge Bemerkung zu machen wie im allgemeinen Fall: die Gleichung $A = 0$ kann eine bedeutungsvolle Aussage sein, wenn nämlich die Definition der Menge A ihr Verschwinden nicht unmittelbar erkennen läßt. Der Fermatsche Satz behauptet: die Menge der natürlichen Zahlen $n > 2$, für welche die Gleichung $x^n + y^n = z^n$ in natürlichen Zahlen x, y, z lösbar ist, ist die Nullmenge.“

Hier wird wieder die Vorstellung über die Welt der Mengen ausgedrückt: Wir können Mengen definieren und mit ihnen operieren, ohne genau über ihren Umfang, ihre Extension, Bescheid zu wissen. Die Menge A aller $n \in \mathbb{N}$, $n > 2$, für die $x^n + y^n = z^n$ in natürlichen Zahlen x, y, z lösbar ist, ist wohldefiniert. A existiert. (Ende des 20. Jahrhunderts hat Andrew Wiles das Fermatsche Problem gelöst: Es gilt $A = \emptyset$.)

Wir kommen nun allgemeiner zu Mengenbildungen durch definierende Eigenschaften.

Mengenbildung über Eigenschaften

Oft will man Objekte mit einer bestimmten Eigenschaft zu einer Menge zusammenfassen, z. B. die Menge der ungeraden natürlichen Zahlen bilden. Hier lautet die zugehörige Eigenschaft $\mathcal{E}(x) =$ „ $x \in \mathbb{N}$ und x ist ungerade“.

Definition

Sei $\mathcal{E}(x)$ eine Eigenschaft. Wir setzen:

$\{x \mid \mathcal{E}(x)\} =$ „die Menge aller Objekte x , auf die $\mathcal{E}(x)$ zutrifft“.

[Die Menge $\{x \mid \mathcal{E}(x)\}$ wird gelesen als: „Menge aller x mit $\mathcal{E}(x)$ “.]

Statt „ $\mathcal{E}(x)$ trifft auf x zu“ sagen und schreiben wir kurz „ $\mathcal{E}(x)$ “.

Es gilt also für alle z :

$$z \in \{x \mid \mathcal{E}(x)\} \text{ gdw } \mathcal{E}(z).$$

Eine genaue Definition von „Eigenschaft“ geben wir im dritten Abschnitt. Hier genügt uns die Intuition: Eine Eigenschaft $\mathcal{E}$ ist eine mathematische Aussage über mathematische Objekte. Es gilt dann für jedes Objekt z : $\mathcal{E}(z)$ oder $\text{non}(\mathcal{E}(z))$.

Zuweilen gefährlich aber suggestiv sind Intelligenztest-Schreibweisen der Form $U = \{1, 3, 5, 7, \dots\}$ für $U = \{x \mid x \in \mathbb{N} \text{ und } x \text{ ist ungerade}\}$.

Oft gibt man Mengen $\{x \mid \mathcal{E}(x)\}$ auch in Form der Aussonderung von bestimmten Elementen aus einer anderen Menge an, z. B.

$$U = \{x \in \mathbb{N} \mid x \text{ ist ungerade}\}.$$

Allgemein:

Definition (*Aussonderung*)

Sei M eine Menge und $\mathcal{E}(x)$ eine Eigenschaft. Wir setzen:

$$\{x \in M \mid \mathcal{E}(x)\} = \{x \mid x \in M \text{ und } \mathcal{E}(x)\}.$$

Wir sammeln hier alle Objekte x mit $\mathcal{E}'(x)$ auf, wobei $\mathcal{E}'(x) =$ „ $x \in M$ und $\mathcal{E}(x)$ “.

Wir betonen schon hier diese Form der Mengenbildung, die *aus einer gegebenen Menge* bestimmte Elemente *aussondert*, da sie auch in der axiomatischen Mengenlehre stets legitim ist, während die unbeschränkte Form $M = \{x \mid \mathcal{E}(x)\}$ bei genauerer Inspektion zu Widersprüchen führt (siehe Kapitel 13).

Die Eigenschaft $\mathcal{E}$ darf fixierte Objekte enthalten:

Definition (*Parameter einer Eigenschaft*)

Sei $\mathcal{E}$ eine Eigenschaft.

Die *Parameter von* $\mathcal{E}$ sind die in $\mathcal{E}$ vorkommenden mathematischen Objekte.

So ist z. B. in $\mathcal{E}(x) =$ „ $x \in \mathbb{N}$ und x ungerade“ die Menge $\mathbb{N}$ ein Parameter und x eine „Variable“. In $M = \{x \mid x \text{ ist kleiner als } 11 \text{ und } x \in U\}$ werden die natürliche Zahl 11 und die kontextabhängige Menge U (hier: der ungeraden Zahlen) als Parameter verwendet.

Jede Menge der Form $\{a_1, \dots, a_n\}$ können wir auch mittels „ $\vee$ “ schreiben, nämlich als

$$\{a_1, \dots, a_n\} = \{x \mid x = a_1 \text{ oder } \dots \text{ oder } x = a_n\}.$$

Hierbei sind dann $a_1, \dots, a_n$ Parameter.

Das naive Komprehensionsprinzip

Lesen wir die Cantorsche Mengendefinition in dem Sinne unvorsichtig, daß sie uns beliebige Zusammenfassungen zu einem Ganzen erlaubt, so können wir das folgende Komprehensionsprinzip für unseren Objektbegriff ableiten:

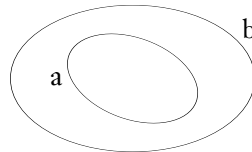
Naives Komprehensionsprinzip für Eigenschaften

Ist $\mathcal{E}(x)$ eine Eigenschaft von mathematischen Objekten, so existiert die Menge $\{x \mid \mathcal{E}(x)\}$ aller Objekte x , auf die $\mathcal{E}(x)$ zutrifft.

Das naive Komprehensionsprinzip ist aber nicht haltbar, es führt zu Widersprüchen. Diese wesentliche Entdeckung besprechen wir im letzten Kapitel der Einführung. Cantor war, wie wir aus Briefen und verschiedenen Bemerkungen in seinen Arbeiten wissen, bereits sehr früh aufgefallen, daß manche sehr große Vielheiten nicht zu Mengen zusammengefaßt werden dürfen; leider hat er aber diese wichtige Erkenntnis nicht besonders betont, und die Lösung der damit verbundenen Schwierigkeiten blieb der nächsten Generation vorbehalten.

Teilmengen

Die neben „ $a \in b$ “ wichtigste Relation zwischen Mengen a und b ist die Teilmengen-Relation.



Definition (Teilmenge und Obermenge)

$$a \subseteq b, \quad b \supseteq a$$

Seien a und b zwei Mengen.

(i) a ist *Teilmenge von* b , in Zeichen $a \subseteq b$, falls gilt:

Für alle $x \in a$ gilt $x \in b$.

(ii) a ist *echte Teilmenge von* b , in Zeichen $a \subset b$, falls $a \subseteq b$ und $a \neq b$.

(iii) Ist $a \subseteq b$, so heißt b *Obermenge von* a , in Zeichen $b \supseteq a$.

Ist $a \subset b$, so heißt b *echte Obermenge von* a , in Zeichen $b \supset a$.

Cantor gebraucht den Ausdruck Teilmenge ab 1884. 1895 definiert er ihn dann gleich im Anschluß an seine „Zusammenfassungs“-Definition (allerdings versteht er unter Teilmengen nur echte Teilmengen, was sich als unpraktische Einschränkung herausstellt).

Das Teilmengensymbol wurde von Ernst Schröder (1841 – 1902) im Jahre 1890 innerhalb seiner „Algebra der Logik“ als „Subsumption zwischen Subjekt und Prädikat“ in der das Eurozeichen vorwegnehmenden Form $\in$ eingeführt. Dieses Zeichen wurde später von der Mengenlehre für die Teilmengenrelation verwendet und zu $\subseteq$ stilisiert. Die Bei-

spiele von Schröder sind von der Form „Gold $\subset$ Metall“. Peano verwendete ein umgedrehtes $\subseteq$ für $\subseteq$, was wieder richtig gedreht wurde, und so wohl Grund dafür wurde, daß in manchen Texten $\subset$, $\not\subset$ anstelle von $\subseteq$, $\not\subseteq$ verwendet wird.

Beispiele: $\{1, 3\} \subseteq \{1, 2, 3\}$, $\{\} \subseteq \{1\}$, $\text{non}(\{1, \{\}\} \subseteq \{1, 2, 3\})$.

Übung

$\subseteq$ ist transitiv: Seien a, b, c Mengen mit $a \subseteq b$ und $b \subseteq c$. Dann gilt $a \subseteq c$.

Statt „ $a \subseteq b$ und $b \subseteq c$ “ schreiben wir auch kürzer „ $a \subseteq b \subseteq c$ “.

Für die Grundobjekte haben wir nach obiger Konvention $\mathbb{N} \subseteq \mathbb{Z} \subseteq \mathbb{Q} \subseteq \mathbb{R}$.

Für alle Mengen M gilt $M \subseteq M$ und $\emptyset \subseteq M$. Für letzteres ist zu überprüfen, ob jedes $x \in \emptyset$ auch in M ist. Es gibt aber gar keine x in $\emptyset$. Also ist wie gewünscht jedes $x \in \emptyset$ in M .

Hausdorff (1914): „Wenn alle Elemente der Menge A auch Elemente der Menge B sind, so sagen wir: A ist in B enthalten, A ist eine Teilmenge von B , eine Menge unter B , oder B enthält A , B ist eine Menge über A . Wir bringen dies durch eine der beiden Formeln

$$A \subseteq B \quad \text{oder} \quad B \supseteq A$$

zum Ausdruck; wobei die Zeichen $\subset$ und $\supset$ an die üblichen Zeichen $<$ und $>$ für kleiner und größer erinnern, aber doch von ihnen unterschieden werden sollen. Zu den Teilmengen von B rechnen wir auch die Menge B selbst und die Nullmenge: eine wichtige Verabredung, deren Zweckmäßigkeit sich vielfach bewähren wird.

Die Teilmengen der aus 4 Elementen bestehenden Menge $\{a, b, c, d\}$ sind:

$$\begin{array}{c} \emptyset \\ \{a\} \quad \{b\} \quad \{c\} \quad \{d\} \\ \{a, b\} \quad \{a, c\} \quad \{a, d\} \quad \{b, c\} \quad \{b, d\} \quad \{c, d\} \\ \{a, b, c\} \quad \{a, b, d\} \quad \{a, c, d\} \quad \{b, c, d\} \\ \{a, b, c, d\} \end{array}$$

Ihre Anzahl ist $1 + 4 + 6 + 4 + 1 = 16 = 2^4$.“

Das Extensionalitätsprinzip können wir nun auch so ausdrücken:

Gleichheitskriterium

Für alle Mengen a, b gilt: $a = b$ gdw $a \subseteq b$ und $b \subseteq a$.

Das Gleichheitskriterium in dieser Form vereinfacht in der Praxis sehr oft den Beweis einer Behauptung $a = b$ für zwei Mengen a, b . In einem ersten Schritt zeigt man $a \subseteq b$, danach zeigt man $b \subseteq a$, und zusammengenommen ergibt sich somit $a = b$.

Einfache Operationen mit Mengen

(Öde für Leser und Autor ist die Einführung der Schnitt- und Vereinigungs- und ähnlicher Dinge. Die elementaren Lehrbücher sind voll von Listen von Gleichungen in größter Allgemeinheit. Hierauf wollen wir ganz verzichten, und hinsichtlich der Beweise solcher Gleichungen raten wir dem Leser, bewehrt mit Papier und Bleistift, sich anhand der bekannten Mengen-Diagramme bestehend aus sich überlappenden Kreisen von der Richtigkeit der Identitäten zu überzeugen. In den folgenden Kapiteln wird es wesentlich spannender ...)

Beim Umgang mit Mengen tauchen die Operationen der Vereinigung, des Durchschnitts und der Subtraktion (oder Differenzbildung) zweier Mengen besonders häufig auf.

Definition (*Vereinigung, Schnitt, Differenz und Disjunktheit*)

Seien a und b zwei Mengen. Die *Vereinigung* $a \cup b$ [a vereinigt b], der *Schnitt* $a \cap b$ [a geschnitten b] und die *Differenz* $a - b$ [a minus b oder a ohne b] von a und b sind definiert durch:

$$a \cup b = \{x \mid x \in a \text{ oder } x \in b\}.$$

$$a \cap b = \{x \mid x \in a \text{ und } x \in b\}.$$

$$a - b = \{x \mid x \in a \text{ und } x \notin b\} = \{x \in a \mid x \notin b\}.$$

Zwei Mengen a, b heißen *disjunkt*, falls $a \cap b = \emptyset$.

Die Symbole „lateinischer Klarheit“ $\cap$ und $\cup$ tauchen zuerst auf in einer Arbeit von Peano aus dem Jahr 1888, diesmal in Italienisch verfaßt:

Peano (1888): „2. Colla scrittura $A \cap B \cap C \dots$, ovvero $ABC \dots$, intenderemo la massima classe contenuta nelle classi $A, B, C, \dots$ ossia la classe formata da tutti gli enti che sono ad un tempo A e B e C , ecc. Il segno $\cap$ si leggerà *e* ...

3. Colla scrittura $A \cup B \cup C \dots$, intenderemo la minima classe che contiene le classi $A, B, C, \dots$, ossia la classe formata dagli enti che sono o A o B o C , ecc. Il segno $\cup$ si leggerà *o* ...“

Das Algebrabuch von Bartel van der Waerden 1930 hat die Zeichenreihe $\in, \subseteq, \cap, \cup$ populär gemacht. Zudem verwendete man lange Zeit Zeichen wie $\mathfrak{G}, \mathfrak{D}$ für den Schnitt und $\mathfrak{S}, \mathfrak{M}$ für die Vereinigung.

Hausdorff (1914): „ A und B seien zwei beliebige Mengen. Unter ihrer Summe

$$S = \mathfrak{S}(A, B)$$

verstehen wir die Menge der Elemente, die mindestens einer der beiden Mengen angehören; unter ihrem Durchschnitt

$$D = \mathfrak{D}(A, B)$$

die Menge der Elemente, die beiden Mengen zugleich angehören.“

$$\begin{array}{ll}
\text{Beispiele:} & \{1, 2\} \cup \{1, 3\} = \{1, 2, 3\}, & \{\} \cup \{1\} = \{1\}, \\
& \{1, 2\} \cap \{1, 3\} = \{1\}, & \{\} \cap \{1\} = \{\}, \\
& \{1, 2\} - \{1, 3\} = \{2\}, & \{\} - \{1\} = \{\}.
\end{array}$$

Für die Subtraktion $a - b$ ist $b \subseteq a$ nicht vorausgesetzt. Oft findet man auch die Schreibweise $a \setminus b$ für $a - b$. (Die englische Lesart „a take away b“ beschreibt sehr gut, was passiert.)

Um die Disjunktheit oder eine disjunkte Vereinigung auszudrücken, stehen einige sprachliche Tricks zur Verfügung, etwa „a zerfällt in b und c“, „a spaltet sich in b und c“ oder „b und c zerlegen a“ für „ $a = b \cup c$ und $b \cap c = \emptyset$ “. Ähnliches gilt für Zerfällungen/Spaltungen/Zerlegungen einer Menge a in mehrere paarweise disjunkte Mengen $a_1, \dots, a_n$, d. h. $a = (\dots((a_1 \cup a_2) \cup a_3) \cup \dots \cup a_{n-1}) \cup a_n$, $a_i \cap a_j = \emptyset$ für alle $i \neq j$, $1 \leq i, j \leq n$. Häufig wird man hier auch fordern, daß keines der a_i die leere Menge ist.

Übung

Seien a, b, c Mengen. Dann gilt:

- (i) $a - b = a - (a \cap b)$,
- (ii) $a - b = a$ gdw $a \cap b = \emptyset$,
- (iii) $a - b = \emptyset$ gdw $a \subseteq b$,
- (iv) $a - (b - c) = (a - b) \cup (a \cap c)$,
- (v) $(a - b) - c = a - (b \cup c)$.

Zur Veranschaulichung sind Diagramme hilfreich; zum (strengeren) Beweis kann man sich am Beweis von (iii) im Satz unten orientieren.

Übung (Assoziativgesetze für Vereinigung und Durchschnitt)

Seien a, b, c Mengen. Dann gilt:

- (i) $(a \cup b) \cup c = a \cup (b \cup c)$,
- (ii) $(a \cap b) \cap c = a \cap (b \cap c)$.

Wir können also Klammern oft weglassen. So schreibt sich etwa

$$a = (\dots((a_1 \cup a_2) \cup a_3) \cup \dots \cup a_{n-1}) \cup a_n$$

viel übersichtlicher als $a = a_1 \cup \dots \cup a_n$. Dagegen ist die Differenzbildung nicht assoziativ, wie (iv) und (v) der vorangehenden Übung zeigen.

Definition (relative Komplemente)

Seien a, b Mengen und $a \subseteq b$.

Dann heißt $b - a$ das relative Komplement von a in b.

Ist b fixiert, so nennen wir $b - a$ kurz das Komplement von a, und setzen $a^c = b - a$.

Satz (*Eigenschaften der relativen Komplemente*)

Sei d eine Menge. Weiter seien $a, b \subseteq d$.

Dann gilt für die relativen Komplemente in d :

- (i) $a \cup a^c = d$,
- (ii) $a \cap a^c = \emptyset$,
- (iii) $(a \cap b)^c = a^c \cup b^c$,
- (iv) $(a \cup b)^c = a^c \cap b^c$.

Die beiden letzten Aussagen werden auch als *de Morgansche Regeln* bezeichnet.

Beweis

zu (i) und (ii): Nach Definition von $a^c = d - a$.

zu (iii): (*Beweistechnik nach dem bekannten Motto:*

If it's madness, there is some method in it.)

zu $\subseteq$: Sei $x \in (a \cap b)^c$.

Also gilt (1) $x \in d$ und (2) $x \notin a$ oder $x \notin b$.

Im ersten Fall von (2) $x \notin a$ ist wegen (1) $x \in d - a = a^c$.

Im zweiten Fall von (2) $x \notin b$ ist wegen (1) $x \in d - b = b^c$.

Also gilt immer: $x \in a^c$ oder $x \in b^c$, also $x \in a^c \cup b^c$.

zu $\supseteq$: Sei $x \in a^c \cup b^c$. Dann gilt $x \in d - a$ oder $x \in d - b$.

Im ersten Fall $x \in d - a$ ist $x \in d - (a \cap b)$ wegen $d - a \subseteq d - (a \cap b)$.

Im zweiten Fall $x \in d - b$ ist $x \in d - (a \cap b)$ wegen $d - b \subseteq d - (a \cap b)$.

Also in beiden Fällen $x \in d - (a \cap b) = (a \cap b)^c$.

Also gilt $(a \cap b)^c \subseteq a^c \cup b^c$ und $(a \cap b)^c \supseteq a^c \cup b^c$.

Nach dem Gleichheitskriterium folgt die Behauptung.

– zu (iv): Analog zu (iii).

Spät, aber nicht ungelegen kommen an dieser Stelle die Distributivgesetze.

Übung (*Distributivgesetze*)

Für alle Mengen a, b, c gilt:

- (i) $(a \cup b) \cap c = (a \cap c) \cup (b \cap c)$,
- (ii) $(a \cap b) \cup c = (a \cup c) \cap (b \cup c)$.

Allgemeinere Formeln zu finden, etwa für $(a \cup b) \cap (c \cup d)$, sei dem Leser als eine Art „open end“-Übung überlassen.

Dem Leser ist vielleicht die Symmetrie zwischen (i) und (ii) aus dem Satz oben, den de Morganschen Regeln und den beiden Distributivgesetzen aufgefallen. Anstelle einer umständlichen Diskussion zitieren wir zur Auflockerung des in dieser Umgebung begrenzt aufregenden Stoffs noch einmal Hausdorff, nämlich über die *Dualität* von $\cup/\cap$, alles/nichts und $\subseteq/\supseteq$. Hierbei ist $\mathfrak{S}(A_1, A_2, \dots, A_m) = A_1 \cup \dots \cup A_m$, $\mathfrak{D}(A_1, \dots, A_m) = A_1 \cap \dots \cap A_m$, und $\bar{A} = A^c = M - A$. Das Zeichen „+“ steht für die disjunkte Vereinigung.

Hausdorff (1914): „Sind $A_1, \dots, A_m$ Teilmengen einer Menge M und $\bar{A}_1, \dots, \bar{A}_m$ ihre Komplemente in M , also

$$M = A_i + \bar{A}_i$$

so ist

$$\begin{aligned} M &= \mathfrak{S}(A_1, A_2, \dots, A_m) + \mathfrak{D}(\bar{A}_1, \dots, \bar{A}_m) \\ &= \mathfrak{D}(A_1, A_2, \dots, A_m) + \mathfrak{S}(\bar{A}_1, \dots, \bar{A}_m); \end{aligned}$$

denn die Elemente von M gehören entweder mindestens einem A_i oder keinem, d. h. entweder der Summe der A_i oder dem Durchschnitt der $\bar{A}_i$ an, und das gleiche gilt, wenn man die A_i mit den $\bar{A}_i$ vertauscht. Wir können diese Formeln kurz so aussprechen: das Komplement einer Summe ist der Durchschnitt der Komplemente, das Komplement eines Durchschnitts die Summe der Komplemente. Ist also P eine Menge, die aus den Mengen A_i durch wiederholte Summen- und Durchschnittsbildung entsteht, so erhält man ihr Komplement $\bar{P}$, indem man die A_i durch ihre Komplemente $\bar{A}_i$, das Zeichen $\mathfrak{S}$ durch $\mathfrak{D}$ und das Zeichen $\mathfrak{D}$ durch $\mathfrak{S}$ ersetzt. Da ferner aus $P = Q$, $P \subset Q$, $P \supset Q$ resp. $\bar{P} = \bar{Q}$, $\bar{P} \supset \bar{Q}$, $\bar{P} \subset \bar{Q}$ folgt, so bleibt jede Gleichung zwischen Mengen richtig, wenn man alle Mengen durch ihre Komplemente ersetzt und die Zeichen $\mathfrak{S}$ und $\mathfrak{D}$ vertauscht; jede Ungleichung bleibt richtig, wenn man außerdem noch die Zeichen $\subset$ und $\supset$ vertauscht [Dualitätsprinzip]. Eine identische, d. h. für beliebige Mengen richtige Relation liefert eine zweite solche auch ohne Übergang zu den Komplementen, also durch bloße Vertauschung der Zeichen $\mathfrak{S}$, $\mathfrak{D}$ und eventuell der beiden Ungleichheitszeichen. Z. B. folgt auf diese Weise die zweite Formel des assoziativen oder distributiven Gesetzes unmittelbar aus der ersten. Als Beispiel für eine Ungleichung zitieren wir die einfache $A \subseteq \mathfrak{S}(A, B)$, aus der durch Dualität $A \supseteq \mathfrak{D}(A, B)$ folgt.“

Weitere einfache Operationen mit Mengen und zugehörige Gleichungen, die gelegentlich in der Mengenlehre und anderswo auftauchen, sind die Themen der folgenden beiden Übungen. Hierzu:

Definition (*symmetrische Differenz*)

Seien a, b Mengen. Dann ist die *symmetrische Differenz* von a und b , in Zeichen $a \Delta b$, definiert durch:

$$a \Delta b = (a - b) \cup (b - a).$$

Ist $\cap$ ein „und“, $\cup$ ein „oder“, so ist Δ ein „entweder oder“.

Übung (*symmetrische Differenz*)

Für alle Mengen a, b, c gilt:

- (i) $a \Delta b = b \Delta a = (a \cup b) - (a \cap b)$,
- (ii) $a \Delta (b \Delta c) = (a \Delta b) \Delta c$,
- (iii) $(a \Delta b) \cap c = (a \cap c) \Delta (b \cap c)$.

In der nächsten Übung betrachten wir geschachtelte Anwendungen der Differenzbildung. Wir vereinbaren zur Vereinfachung der Notation Rechtsklammerung, d. h. $a - b - c = a - (b - c)$, $a - b - c - d = a - (b - c - d)$, usw.

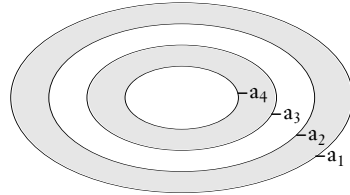
Differenzenketten sind von Hausdorff untersucht worden und spielen in der deskriptiven Mengenlehre eine Rolle. Hier sind sie lediglich ein hübsches Übungsmaterial, das viel attraktiver ist als etwa die Distributivgesetze:

Übung (Differenzenketten aus absteigenden Mengen)

Sei $n \in \mathbb{N}$, und seien $a_1, \dots, a_n$ Mengen mit $a_1 \supseteq a_2 \supseteq \dots \supseteq a_n$.

Dann gilt:

- (i) Für n gerade: $a_1 - \dots - a_n = (a_1 - a_2) \cup \dots \cup (a_{n-1} - a_n)$.
- (ii) Für n ungerade: $a_1 - \dots - a_n = (a_1 - a_2) \cup \dots \cup (a_{n-2} - a_{n-1}) \cup a_n$.



Differenzenkette der Länge 4.

Der Leser kann sich also Differenzenketten beliebiger Länge visualisieren: Man wirft einen Stein ins Wasser und sammelt, je nachdem ob die Länge n der Kette gerade ist oder ungerade, die Ringe bestehend aus Wellenbergen oder Wellentälern.

Mengenbildung über Eigenschaften und Operationen

Oft will man Mengenbildungen der folgenden Art durchführen: x durchläuft alle Elemente einer Menge M und wird dabei durch eine bestimmte Operation zu einem neuen Objekt y umgewandelt; alle so erhaltenen Objekte y sollen zu einer Menge N zusammengefaßt werden. Etwa könnte $y = \{x\}$ sein, und wir wollen dann die Menge N aller $\{x\}$ mit $x \in M$ bilden. Es ist sehr suggestiv, dies in der folgenden Form zu notieren: $N = \{\{x\} \mid x \in M\}$. Diese Schreibweise wollen wir nun etwas präzisieren.

Definition

Sei $\mathcal{E}(x)$ eine Eigenschaft, und sei $\mathcal{F}(x)$ eine Operation. Wir setzen:

$$\{\mathcal{F}(x) \mid \mathcal{E}(x)\} = \{y \mid \text{es gibt ein } x \text{ mit } \mathcal{E}(x) \text{ und } y = \mathcal{F}(x)\}.$$

[$\{\mathcal{F}(x) \mid \mathcal{E}(x)\}$ wird gelesen als: „Menge aller $\mathcal{F}(x)$ mit $\mathcal{E}(x)$ “.]

In dieser Definition haben wir die neue Form $\{\mathcal{F}(x) \mid \mathcal{E}(x)\}$ auf die alte Form $\{y \mid \mathcal{E}'(y)\}$ zurückgeführt. Es gilt:

$$\{\mathcal{F}(x) \mid \mathcal{E}(x)\} = \text{„die Menge aller Objekte } \mathcal{F}(x), \text{ auf deren Argument } x \text{ die Eigenschaft } \mathcal{E}(x) \text{ zutrifft“}.$$

In dieser Form wird $\{\mathcal{F}(x) \mid \mathcal{E}(x)\}$ intuitiv auch verstanden: Wir sammeln alle $\mathcal{F}(x)$ auf, wobei x bestimmte Bedingungen erfüllt.

Ebenso wie wir nicht genau definiert haben, was eine Eigenschaft $\mathcal{E}(x)$ ist, so haben wir hier nicht definiert, was eine Operation $\mathcal{F}(x)$ ist. Intuitiv ist eine Operation eine Zuordnung von Objekten. Einem Objekt x wird in eindeutiger Weise

durch $\mathcal{F}$ ein Objekt y zugeordnet, und dieses wird als $\mathcal{F}(x)$ bezeichnet. In konkreten Fällen läßt sich aber die Zusammenfassung aller $\mathcal{F}(x)$ mit der Nebenbedingung $\mathcal{E}(x)$ nicht nur auf die alte Form zurückführen, sondern auch der Operationsbegriff kann dabei eliminiert werden. So ist etwa $\{ \{ x \} \mid x \in M \}$ nach Definition identisch mit $\{ y \mid \text{es gibt ein } x \in M \text{ mit } y = \{ x \} \}$, und diese Menge hätten wir bereits vor der obigen Definition problemlos bilden können. Kurz: Die Mengenbildung über Eigenschaften *und* Operationen kann man als eine bequeme neue Schreibweise für eine Mengenbildung über Eigenschaften auffassen. So gesellt sich zum etwas vagen Eigenschaftsbegriff keine neue Ungenauigkeit hinzu.

Diese ausführliche Diskussion mag dem Leser vielleicht etwas pedantisch erscheinen, und er hätte sicher $\{ \{ x \} \mid x \in M \}$ ohne weitere Erläuterung verstanden. Sie wird aber gerechtfertigt durch die Tatsache, daß in der axiomatischen Mengenlehre, wo die uneingeschränkte Komprehension $\{ x \mid \mathcal{E}(x) \}$ nicht mehr zur Verfügung steht, die Mengenbildung $\{ \mathcal{F}(x) \mid x \in M \}$ für eine Menge M und eine Operation $\mathcal{F}$ durch ein eigenes, recht starkes Axiom gefordert werden muß, und daß zudem dieses auf Abraham Fraenkel (1922) u. a. zurückgehende *Ersetzungsaxiom* viele Jahre nach der Einführung der Axiomatik von Ernst Zermelo 1908 nicht beachtet wurde. Die Bildung von $\{ \mathcal{F}(x) \mid x \in M \}$ bringt intuitiv zusätzliche Dynamik und Komplexität in den Akt der Zusammenfassung, unbeschadet der Tatsache, daß sie auf eine einfache Komprehension zurückgeführt werden kann.

Die Operation $\mathcal{F}$ kann auch mehrstellig sein, und je n Objekten $x_1, \dots, x_n$ ein Objekt $\mathcal{F}(x_1, \dots, x_n)$ zuordnen. Sie darf wie eine Eigenschaft Parameter enthalten. Häufig ist eine Operation auch nicht auf allen Objekten definiert, sondern nur auf den Mengen oder auch nur auf den Elementen einer bestimmten Menge.

Beispiele für Operationen sind etwa:

$$\mathcal{F}(x) = \{ x \},$$

$$\mathcal{F}(x) = x \cup a \quad \text{für Mengen } x \text{ mit einer festen Menge } a \text{ als Parameter,}$$

$$\mathcal{F}(x, y) = (x \Delta y \Delta a) \cap b \quad \text{für Mengen } x, y \text{ und Mengenparametern } a, b,$$

$$\mathcal{F}(x_1, \dots, x_5) = (x_1 - (x_2 \cup x_3)) \cap x_4 \cap x_5 \quad \text{für Mengen } x_1, \dots, x_5.$$

Mengensysteme

Wir brauchen noch allgemeinere Versionen des Durchschnitts und der Vereinigung. Diese werden für Mengensysteme definiert:

Definition (*Mengensysteme*)

Ein *Mengensystem* M ist eine Menge, deren Elemente alle Mengen sind, d. h. M enthält keine Grundobjekte als Elemente.

Definition (*Großer Durchschnitt und große Vereinigung*)

Sei M ein Mengensystem. Dann sind der *Durchschnitt von M* , in Zeichen $\bigcap M$, und die *Vereinigung von M* , in Zeichen $\bigcup M$, wie folgt definiert:

$$\bigcap M = \{x \mid \text{für alle } a \in M \text{ ist } x \in a\},$$

$$\bigcup M = \{x \mid \text{es existiert ein } a \in M \text{ mit } x \in a\}.$$

Die Vereinigung $\bigcup M$ und der Durchschnitt $\bigcap M$ sind weitere Beispiele für einstellige Operationen $\mathcal{F}$ auf Mengen.

Beispiele: Für alle Mengen a, b, c gilt

$$\bigcap \{a, b\} = a \cap b,$$

$$\bigcup \{a, b, c\} = a \cup b \cup c,$$

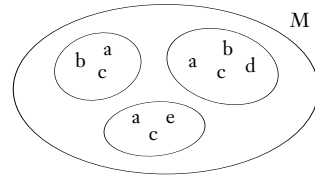
$$\bigcap \{a\} = \bigcup \{a\} = a.$$

Streng nach Definition gilt $\bigcup \emptyset = \emptyset$ und

$\bigcap \emptyset = \{x \mid x \text{ ist Objekt}\}$. Die erste Aussage ist klar. Zum Beweis der zweiten

Aussage sei x beliebig. Wir zeigen $x \in \bigcap \emptyset$. Hierzu ist zu zeigen: Für alle $a \in \emptyset$

gilt $x \in a$. Es gibt aber gar keine $a \in \emptyset$, also ist die Aussage sicher richtig.



$$\bigcap M = \{a, c\}, \quad \bigcup M = \{a, b, c, d, e\}$$

Übung

Es gilt $\bigcap \{\{m \in \mathbb{N} \mid m \geq n\} \mid n \in \mathbb{N}\} = \emptyset$.

Erfahrungsgemäß ist der Umgang mit großen Vereinigungen und Schnitten und die Rolle der leeren Menge bei Erstkontakt etwas unangenehm. Diese Dinge werden aber mit der Zeit trivial.

Wir kommen nun zu einer harmlos aussehenden Operation, die zu den interessantesten der Mengenlehre gehört, weil sie für unendliche Mengen schlecht verstanden ist – wie wir sehen werden.

Die Potenzmenge

Eine entscheidende Rolle bei der Untersuchung der Größe einer Menge wird die Potenzmengenoperation spielen:

Definition (*Potenzmenge*)

Sei M eine Menge. Dann ist die *Potenzmenge* von M , in Zeichen $\mathcal{P}(M)$, die Menge aller Teilmengen von M :

$$\mathcal{P}(M) = \{a \mid a \subseteq M\}.$$

Die Potenzmenge rückt seit Zermelo in den Mittelpunkt des Interesses. 1908 führt er den Begriff ein und schreibt 2^M für die Potenzmenge einer Menge, wobei das 2 an „Unter-mengen“ erinnert. Der Name Potenzmenge bietet sich wegen des Zusammenhangs mit der arithmetischen Potenzoperation an (vgl. die Übung unten). Gerhard Hessenberg (1874 – 1925) spricht in seinem Lehrbuch von 1906 noch von der „Menge der Teilmengen“, ohne einen kompakten Begriff zu gebrauchen.

Beispiele: $\mathcal{P}(\emptyset) = \{\emptyset\}$, $\mathcal{P}(\{x\}) = \{\emptyset, \{x\}\}$,
 $\mathcal{P}(\{\emptyset, \{\emptyset\}\}) = \{\emptyset, \{\emptyset\}, \{\{\emptyset\}\}, \{\emptyset, \{\emptyset\}\}\}.$

Für alle Mengen M gilt $\emptyset \in \mathcal{P}(M)$ und $M \in \mathcal{P}(M)$. Ist weiter $a \in \mathcal{P}(M)$, so ist auch $M - a \in \mathcal{P}(M)$. $\mathcal{P}(M)$ ist darüber hinaus abgeschlossen unter Schnitten und Vereinigungen: Ist $N \subseteq \mathcal{P}(M)$, so sind $\bigcap N$ und $\bigcup N$ Elemente von $\mathcal{P}(M)$.

Übung

- (i) Man bestimme die Potenzmenge von $\{1\}$, $\{1, 2\}$, $\{1, 2, 3\}$, und zähle ihre Elemente.
- (ii) Wieviele Elemente hat $\mathcal{P}(\{1, \dots, n\})$ für $n \in \mathbb{N}$?
- (iii) Wieviele Teilmengen von $\{1, \dots, n\}$ mit genau k Elementen gibt es für $n \in \mathbb{N}$ und $0 \leq k \leq n$?

Obwohl mit großer Vorsicht zu genießen, beenden wir dieses Kapitel mit einer Anekdote von Felix Bernstein (1878 – 1956), die uns überliefert, wie sich Cantor und Dedekind Mengen vorgestellt haben. Sie wird wahrscheinlich auch deswegen immer wieder erzählt, weil sie die großen Unterschiede zwischen den Charakteren Cantor und Dedekind mit wenigen Strichen nachzeichnet. Das Zitat findet sich in den „Gesammelten Werken“ von Richard Dedekind, und wurde dort von der Herausgeberin Emmy Noether innerhalb eines Kommentars eingefügt.

Georg Cantors Vorstellung von Mengen, berichtet von Felix Bernstein

„F. Bernstein übermittelt noch die folgenden Bemerkungen:
 , ... Von besonderem Interesse dürfte folgende Episode sein: Dedekind äußerte, hinsichtlich des Begriffes der Menge: er stelle sich eine Menge vor wie einen geschlossenen Sack, der ganz bestimmte Dinge enthalte, die man aber nicht sähe, und von denen man nichts wisse, außer daß sie vorhanden und bestimmt seien. Einige Zeit später gab Cantor seine Vorstellung einer Menge zu erkennen: Er richtete seine kolossale Figur hoch auf, beschrieb mit erhobenem Arm eine großartige Geste und sagte mit einem ins Unbestimmte gerichteten Blick: ‚Eine Menge stelle ich mir vor wie einen Abgrund.‘“

(in: Dedekind 1930 – 1932, *Gesammelte Werke*, Bd. III, S. 449)

2. Zwischenbetrachtung

Kritik oder „Sturm und Drang“ ?

An dieser Stelle bieten sich zwei Möglichkeiten an fortzufahren:

1. **Weg** Kritische Betrachtung der Begriffe.
2. **Weg** Untersuchung der Intuition auf ihren mathematischen Gehalt.

Der erste Weg

Eine präzisierende Analyse des naiven Verständnisses der Begriffe *Menge*, „ $a \in b$ “, *Mengenbildung*, *Eigenschaft* führt fast zwangsläufig zur axiomatischen Mengenlehre, die in einer ebenso einfachen wie strengen Kunstsprache formuliert ist, der sogenannten Prädikatenlogik erster Stufe.

Dieser Weg könnte etwa wie folgt verlaufen (historisch verlief die Sache auf dem zweiten Weg). Zunächst kann man die folgende Frage stellen:

A. Was ist eigentlich ein mathematisches Objekt?

Es zeigt sich, daß alle mathematischen Objekte (Zahlen, Funktionen, usw.) als Mengen interpretiert werden können. Dies heißt: Es gibt Definitionen dieser Objekte als Mengen, die alle Eigenschaften dieser Objekte zur Verfügung stellen, die in der Mathematik gebraucht werden. Hier geht es nicht um Ontologie – was ist π ? –, sondern um präzise und brauchbare, sich im Aufbau einer Theorie natürlich ergebende Definitionen – z. B. $\pi/2$ ist definiert als „die kleinste positive Nullstelle der Cosinus-Funktion“. Was man sich unter den mathematischen Objekten schließlich vorstellt und welche Eigenschaften der Objekte am wichtigsten erscheinen, ist jedem Mathematiker selbst überlassen – z. B. π als fundamentale Größe zur Berechnung von Umfang und Inhalt des Kreises.

Die Mengenlehre ist hinsichtlich der Interpretation der gesamten Mathematik konkurrenzlos. Entscheidend ist hier nicht ein platonischer Glaube an die Mengen, sondern die Leistungsfähigkeit der Theorie und die Universalität der verwendeten Sprache.

Hat man nun gesehen, daß sich die Mengenlehre als Rahmentheorie für die Mathematik eignet, wird man den Begriff der Menge selbst hinterfragen und Beweise, die ganz in der Sprache des Mengenbegriffs geführt sind, genauer betrachten. Neben logischen Schlüssen findet man hier insbesondere zahllose Existenzbehauptungen, etwa in der Form „eine Menge mit den und den Eigenschaften

existiert“ oder „aus einer oder mehr Mengen können neue Mengen gebildet werden“, zu M, N z. B. $\{M, N\}$ und zu M etwa $\mathcal{P}(M)$. Man wird also fragen:

B. Welche mathematischen Objekte (= Mengen) existieren?

Zur Beantwortung dieser Frage wird man Axiome $\mathfrak{M}$ angeben, die Eigenschaften der Elementbeziehung wie das Extensionalitätsprinzip wiedergeben, und die in sorgfältiger Auswahl die Existenz von bestimmten Mengen garantieren. Sorgfalt ist deswegen nötig, weil man schnell sieht, daß das naive Komprehensionschema zu Widersprüchen führt; es ist leicht, zu viel Existenz zu fordern, und das System dadurch zu ruinieren. Man wird sich bei der Aufstellung der Axiome sowohl an der Intuition über „Menge, Element“ orientieren als auch am mathematischen Bedarf. Ein Axiom von $\mathfrak{M}$ wird z. B. lauten:

Potenzmengenaxiom

Zu jeder Menge M existiert die Potenzmenge von M , d. h.:

Für alle M gibt es ein N mit der Eigenschaft:

Für alle x gilt: $x \in N$ gdw $x \subseteq M$.

$\mathfrak{M}$ ist hier eine Bezeichnung für ein „externes“ Axiomensystem, das die Mengenwelt selber – für Platoniker – beschreibt bzw. – für Formalisten – regelt, nicht für ein Objekt innerhalb der Mengenwelt.

Weiter ist eine Präzisierung des Eigenschaftsbegriffs erforderlich: Welche Ausdrücke $\mathcal{E}$ in der Aussonderung $\{x \in M \mid x \text{ hat die Eigenschaft } \mathcal{E}\}$ sind erlaubt oder möglich? Die genaue Formulierung der Axiome und die Präzisierung des Eigenschaftsbegriffs führen zur Prädikatenlogik erster Stufe. Das Axiom über die Existenz der Potenzmenge schreibt sich darin so:

„ $\forall M \exists N \forall x (x \in N \leftrightarrow \forall y (y \in x \rightarrow y \in M))$ “, gelesen:

„Für alle M existiert ein N , sodaß für alle x gilt:

x ist in N genau dann, wenn für alle y gilt: y in x folgt y in M .“

Mit der Abkürzung $a \subseteq b$ für $\forall c (c \in a \rightarrow c \in b)$ erhält man eine besser lesbare Form des Potenzmengenaxioms, die obiger Formulierung entspricht:

$\forall M \exists N \forall x (x \in N \leftrightarrow x \subseteq M)$.

Solche Abkürzungen entsprechen genau den Definitionen neuer Konzepte – in diesem Fall dem der Teilmengenrelation. Einer kargen Grundsprache wird so Schritt für Schritt ein umfassendes Lexikon an die Seite gestellt, mit dessen Wortschatz man über komplexe Begriffe leicht reden kann, ohne dabei prinzipiell jemals mehr sagen zu können als ganz zu Beginn. Im Gegensatz zur natürlichen Sprache ist dieses Lexikon aber nicht einfach alphabetisch geordnet, sondern es gilt die Verabredung, daß ein neuer Eintrag neben der Verwendung der Grundsprache ausschließlich auf weiter vorne stehende Lexikon-Einträge zu verweisen hat.

Eigenschaften $\mathcal{E}$ und mathematische Aussagen sind dann einfach bestimmte $\forall, \exists$ -Ausdrücke, sogenannte Formeln in der Sprache der Mengenlehre. Was eine Formel ist und was nicht, wird natürlich genau festgelegt.

Mit Hilfe der Axiome kann man nun die Existenz vieler Mengen folgern und mathematische Zusammenhänge zwischen Mengen beweisen. Dies führt zur Frage:

C. Was genau ist ein mathematischer Beweis einer Aussage?

Hier läßt sich ein Kalkül angeben, der aus Regeln besteht, wie bestimmte Zeichenketten unserer Kunstsprache in andere verwandelt werden dürfen. Das Ergebnis einer solchen schrittweisen Umformung liefert einen mathematischen Satz, und die Reihe der Umformungen selber bildet einen formalen Beweis dieses Satzes. Die Beweise, die üblicherweise in der Mathematik in einer reduzierten Umgangssprache, dem mathematischen Deutsch oder Englisch etwa, gegeben werden, lassen sich mit den strengen formalen Manipulationsregeln des Kalküls nachbauen, wenn auch nur in sehr mühevoller Weise.

Was hat man damit erreicht? Ein Axiomensystem und einen präzisen Beweisbegriff für die Mathematik. Jeder irgendwo auf der Welt geführte mathematische Beweis gleich welcher Disziplin läßt sich auf der Basis der Axiome der Mengenlehre durchführen und zudem – zumindest theoretisch – formalisieren, d. h. er kann innerhalb der Kunstsprache formuliert und mechanisch auf seine Richtigkeit überprüft werden. Angestrebt wird keinesfalls die *Ersetzung* der üblichen semantischen (inhaltlichen) mathematischen Kultur durch eine syntaktische (formale), sondern ihre Bereicherung durch das Wissen um eine *prinzipielle Übertragbarkeit* dieser Kultur in einen formalen Rahmen. Nur dadurch werden Fragen und Ergebnisse *über* die Mathematik, etwa: *Was ist beweisbar?* möglich. Die Analyse der Mathematik selbst und die dabei verwendeten Methoden und erzielten Resultate bezeichnet man seit David Hilbert (1862 – 1943) als *Metamathematik*. Wir kommen im dritten Abschnitt darauf zurück.

Der formale Beweisbegriff ermöglicht es daneben auch, für die Beweisfindung, Beweisüberprüfung und Beweisanalyse Computer einzusetzen, die sich bei der syntaktischen Manipulation von weltumspannend langen Zeichenketten wesentlich wohler fühlen als in Gesellschaft mancher dreizeiliger Beweise aus einem Lehrbuch. Das ist alles noch in den Kinderschuhen, auch wenn in den letzten Jahrzehnten zuweilen eine neue Schuhgröße notwendig wurde. Die meisten Mathematiker bezweifeln, daß der Computereinsatz in der Beweisführung je ein ähnliches Niveau erreichen wird wie beim Schachspiel. Es gibt Resultate der mathematischen Logik, die die Grenzen der mechanischen Beweisführung betreffen, und die den Traum von der unersetzbaren „biologischen“ Kreativität in der Mathematik am Leben erhalten.

Hat man nun ein geeignetes Axiomensystem zusammengestellt, so erhebt sich bald die Frage nach seiner Leistungsfähigkeit:

- (L1) Ist das Axiomensystem widerspruchsfrei?
- (L2) Ist das Axiomensystem vollständig?

Die Bejahung von (L1) bedeutet einfach: Die Aussage $0 = 1$ (oder $\exists x x \neq x$) läßt sich nicht formal aus den Axiomen ableiten.

Die Bejahung von (L2) bedeutet: Für jede Aussage ϕ existiert ein Beweis von ϕ oder ein Beweis der Negation von ϕ , in Zeichen $\neg \phi$, gelesen: non ϕ , d. h. jede Aussage läßt sich beweisen oder widerlegen.

Zu diesen beiden Fragen hat Kurt Gödel (1906 – 1978), auf den Schultern von David Hilbert und Bertrand Russell (1872 – 1970), die den Sprachraum für diese Probleme zur Verfügung stellten, in den dreißiger Jahren des 20. Jahrhunderts fundamentale Resultate erzielt – die Gödelschen Unvollständigkeitssätze. Sie beantworten die Frage (L1) mit: „Wir können es nicht sicher wissen: die Widerspruchsfreiheit der axiomatischen Mengenlehre ist unbeweisbar.“ Und die Frage (L2) beantworten sie schlichtweg mit „Nein!“: Ist das zugrunde gelegte Axiomensystem der Mengenlehre widerspruchsfrei, so gibt es stets Aussagen, die weder beweisbar noch widerlegbar sind. (Ist das Axiomensystem widerspruchsvoll, verliert die Frage (L2) ihren Sinn, da dann jede Aussage beweisbar ist.) Wir werden auf die Gödelschen Sätze noch mehrfach zurückkommen. Beweise und Erläuterungen findet der Leser in den Lehrbüchern zur mathematischen Logik.

Eine weitere klassische Frage an ein Axiomensystem ist:

(L3) Ist das Axiomensystem unabhängig?

Die Bejahung von (L3) bedeutet: Ist ϕ ein Axiom des Systems, so läßt sich ϕ nicht aus den übrigen Axiomen des Systems beweisen. Ein Axiomensystem ist also unabhängig, wenn jedes Mitglied des Systems die Stärke des Systems erhöht. Obwohl dieser Frage wesentlich weniger Bedeutung zukommt als den beiden anderen, so führt doch die Klärung der inneren Abhängigkeiten zu einem besseren Gesamtverständnis des Systems.

Eine weitere Frage an ein Axiomensystem wäre die nach besonders einfachen gleichwertigen Systemen, etwa solchen, die nur endlich viele Axiome enthalten. (Zwei Axiomensysteme sind hierbei gleichwertig, wenn sich jedes Axiom des einen Systems im anderen System beweisen läßt und umgekehrt.) Das heute zumeist verwendete Axiomensystem von Zermelo-Fraenkel ZFC für die Mengenlehre hat unendlich viele Mitglieder, und man kann zeigen, daß es kein endliches gleichwertiges System gibt (vgl. hierzu jedoch auch Abschnitt 3, Kapitel 3).

Schließlich sei zur Frage (L1) noch bemerkt, daß man hier doch nicht völlig im dunklen Wald stehen bleiben muß. Es ist gelungen, von bestimmten Axiomen ϕ der Mengenlehre ihre relative Widerspruchsfreiheit nachzuweisen, d. h.: Ist die Axiomatik widerspruchsvoll, so ist bereits die Axiomatik ohne ϕ widerspruchsvoll. Das Axiom ϕ erhält dadurch in gewisser Weise einen Persilschein seiner Widerspruchsfreiheit. Insbesondere ist für das sehr kritisch beäugte Zermelosche Auswahlaxiom ein Nachweis der relativen Widerspruchsfreiheit möglich (Gödel 1938, vgl. hierzu auch die Erläuterungen innerhalb der Zeittafel zur Mengenlehre).

Der zweite Weg

Warum aber soll man überhaupt mit einer kritischen Analyse beginnen? Die Relation „ $a \in b$ “ erscheint zunächst klar, ungefährlich und erweist sich, zusammen mit dem Begriff einer Funktion, als fruchtbar und interessant. Und selbst die klare Erkenntnis der inneren Widersprüche allzu uferloser Zusammenfassungen von Objekten zu Mengen hat Mathematiker wie Cantor und Hausdorff in keiner Weise davon abgehalten, die Mengenlehre nach metaphysisch-ästhetischen Kriterien zu errichten und nach ästhetischen Kriterien weiterzuentwickeln. Hierfür sind vielfach nur gut überschaubare und relativ kleine Zusammenfassungen nötig, und bereits im Umfeld der reellen Zahlen ergeben sich ebenso

schwierige wie fesselnde Fragen, die auf zuweilen störende Rückenprobleme wie ein Betäubungsmittel wirken können, und neben denen Risse im Fundament als tolerierbare Bauungenauigkeiten erscheinen. Dem Mengenbegriff ist darüber hinaus auch nicht so ohne weiteres anzusehen, daß sich auf ihm eine Universal-sprache für die Mathematik gründen läßt. Zunächst liegt also eine „Sturm und Drang“-Periode nahe, in der die Begriffe auf ihren Gehalt und ihren inneren Reichtum untersucht werden. Auch wir werden uns bis zum dritten Abschnitt weiter an diese faustische Mengenlehre halten, in Übereinstimmung mit der historischen Entwicklung.

Der Anfang ist immer das Schwerste bei einer Sache, und Cantor hatte zusätzlich mit Verboten zu kämpfen, mit denen Ende des 19. Jahrhunderts das aktual Unendliche, die „fertige“ unendliche Menge belegt war. Heute – die unendliche Menge der natürlichen Zahlen $\mathbb{N}$ ist jedem Schüler ein Begriff geworden – ist der Anfang der Theorie der unendlichen Mengen aber leicht zu bestreiten, und nicht begründbare Dogmen sind zerbrochen. Ein Stück weit werden wir diesen Weg nun gehen, und uns hierbei auf die Intuition verlassen. Die faszinierenden Phänomene der Größenunterschiede im Unendlichen können so vielleicht am deutlichsten hervortreten. Zur Natürlichkeit der naiven Untersuchung der Begriffe gesellt sich heute zudem die klärende Wirkung historischer Distanz.

Langfristig ist aber die Durchführung der kritischen Analyse unvermeidbar. Die heutige Mengenlehre hat nach dieser – insgesamt mehrere Jahrzehnte dauernden – Durchführung alle wesentlichen Ergebnisse und Konzepte der nicht-kritischen „klassischen“ Phase retten können, ihren Gehalt herausgearbeitet, und sie auf ein solides Fundament gestellt. Der Leser muß also nicht fürchten, nachher alles wieder vergessen oder neu lernen zu müssen, sondern darf sich vielmehr darauf freuen, vom Parkett in die Logen umzuziehen: Da das Theater eine Unzahl interessanter Stücke zu bieten hat, lohnt sich der bessere Blick. Die Stimmung dort unten sollte man aber einmal erlebt haben.

Aus Abraham Fraenkels Einleitung zu „Mengenlehre und Logik“

„... vielmehr sollen diejenigen Grundgedanken der *abstrakten Mengenlehre* möglichst einfach entwickelt werden, die in enger Beziehung zu *logischen* Problemen und Methoden stehen, und eben diese Beziehungen grundsätzlich herausgearbeitet werden. Es trifft sich glücklich für die Bedürfnisse des mathematisch oder logistisch nicht vorgebildeten Lesers, daß das genannte Ziel in weitem Ausmaß ohne nennenswerte mathematische Technik und ohne symbolisch-logische Einkleidung erreichbar ist.“

(Abraham Fraenkel 1959)

3. Abbildungen zwischen Mengen

In diesem Kapitel führen wir den für alles Folgende grundlegenden Begriff der Abbildung oder Funktion rein mengentheoretisch ein: Eine Funktion ist hier nach eine „statische“ Menge mit bestimmten Eigenschaften, kein neuer Grundbegriff. Die Intuition, daß eine Funktion eine aktive Aufgabe wahrnimmt, um ein Objekt in ein anderes überzuführen oder ihm ein anderes zuzuordnen, bleibt davon unberührt.

Geordnete Paare

Für sich nützlich und für den Funktionsbegriff unentbehrlich ist der Begriff des geordneten Paares P zweier Objekte a und b , in Zeichen $P = (a, b)$. Man könnte geordnete Paare als Grundbegriff betrachten, aber wir wollen sie hier auf den Mengenbegriff zurückführen, nicht zuletzt als Beispiel für eine mengentheoretische Interpretation eines mathematischen Konstrukts.

Für Mengen gilt immer $\{a, b\} = \{b, a\}$ nach dem Gleichheitskriterium. Bei dem geordneten Paar von a und b soll dagegen die Reihenfolge der Elemente eine Rolle spielen. Entscheidend ist offenbar die Bedingung:

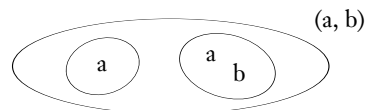
(#) $(a, b) = (c, d) \text{ gdw } a = c \text{ und } b = d$.

Dies ist die einzige Eigenschaft, die wir von einem geordneten Paar erwarten, und wir können irgendeine Definition nehmen, die sie erfüllt. Bequem – und allgemein üblich geworden – ist die folgende Definition von Kazimierz Kuratowski (1896 – 1980) aus dem Jahr 1921:

Definition (*geordnetes Paar; Kazimierz Kuratowski*)

Seien a, b Objekte. Dann ist *das geordnete Paar von a und b* , in Zeichen (a, b) , definiert durch:

$$(a, b) = \{\{a\}, \{a, b\}\}.$$



Übung

Man zeige (#), d. h. für alle Objekte a, b, c, d gilt:

$$(a, b) = (c, d) \text{ gdw } a = c \text{ und } b = d.$$

Das geordnete Paar (a, b) hat zwei Elemente für $a \neq b$, und nur ein Element für $a = b$, nämlich $\{a\}$, denn es gilt $(a, a) = \{\{a\}, \{a, a\}\} = \{\{a\}\}$.

Hausdorff (1914): „Aus zwei verschiedenen Elementen a, b können wir die Menge oder das Paar $\{a, b\} = \{b, a\}$ zusammensetzen; beide Elemente treten darin symmetrisch, gleichberechtigt auf. Wir können sie aber auch zu einer unsymmetrischen, das eine Element vor dem anderen bevorzugenden Verbindung zusammenfassen: wir können das geordnete Paar

$$p = (a, b)$$

bilden, das von dem umgekehrt geordneten

$$p^* = (b, a)$$

unterschieden werden soll. Falls die beiden Elemente gleich sind, können wir sie einerseits zur Menge $\{a\}$, andererseits zu dem geordneten Paar (a, a) zusammenfassen, das in diesem Fall mit seiner Umkehrung identisch ist. Zwei geordnete Paare $p = (a, b)$ und $p' = (a', b')$ gelten dann und nur dann als gleich, wenn $a = a'$ und $b = b'$ [ist].

Die Doppelindizes (i, k) an den Elementen einer Determinante, die rechtwinkligen Koordinaten (x, y) von Punkten der Ebene sind geordnete Zahlenpaare. Dieser Begriff ist also in der Mathematik fundamental, und die Psychologie würde hinzufügen, daß geordnete, unsymmetrische, selektive Verknüpfung zweier Dinge sogar ursprünglicher ist als ungeordnete, symmetrische, kollektive. Denken, Sprechen, Lesen und Schreiben sind an zeitliche Reihenfolge gebunden, die sich uns aufzwingt, bevor wir von ihr absehen können. Das Wort ist früher da als die Menge seiner Buchstaben, das geordnete Paar (a, b) früher als das Paar $\{a, b\}$.

Übrigens läßt sich, wenn man will, der Begriff des geordneten Paares auf den Mengenbegriff zurückführen. Sind 1, 2 zwei voneinander wie von a und b verschiedene Elemente, so hat das Paar von Paaren

$$\{\{a, 1\}, \{b, 2\}\}$$

genau die formalen Eigenschaften des geordneten Paares (a, b) , nämlich die Unvertauschbarkeit von a und b im Falle der Verschiedenheit beider Elemente ... “

Die von Hausdorff unter „übrigens“ erwähnte Paardefinition hat bei aller Natürlichkeit den Nachteil, daß sie zusätzliche Objekte mit ins Spiel bringt. Für das Paar $(1, 2)$ selbst brauchen wir zudem neue Platzindikatoren $1', 2'$. Dieser Indikatorenwechsel in Hausdorffs Definition ist im Prinzip nicht notwendig:

Übung (aus der Juristenabteilung der Mengenlehre)

Seien i, j zwei verschiedene Objekte. Dann gilt (#) für die Paardefinition:

$$(a, b) = \{\{a, i\}, \{b, j\}\} \quad \text{mit beliebigen Objekten } a, b.$$

Die Raffinesse der Kuratowski-Definition ist aber, daß in ihr lediglich die beiden Objekte a, b auftauchen, deren Paar gebildet werden soll, und keine Indikatoren. Nicht alle mengentheoretischen Interpretationen mathematischer Konstrukte sind gleich gut.

Kuratowski (1921): „Nous terminons cette note par une remarque suivante sur la notion de paire ordonnée.

Soit A un ensemble composé de deux éléments a et b .

Il n'existe que deux classe, qui établissent un ordre dans A , à savoir:

$$\{\{a, b\}, \{a\}\} \quad \text{et} \quad \{\{a, b\}, \{b\}\}.$$

Il semble bien naturel d'admettre la définition suivante:

Définition V. La classe $\{\{a, b\}, \{a\}\}$ est une „*paire ordonnée*“ dont a est le premier élément et b est le second‘.

La notion de paire ordonnée est, comme on sait une des plus importantes dans la théorie des ensembles et il est bien utile d’avoir pour elle une définition suffisamment simple. En admettre celle que nous venons de proposer est une conséquence immédiate de l’emploi de la théorie de l’ordre qui a été discutée ici.“

Kuratowski notiert im Original Mengen mit runden Klammern anstatt mit geschweiften. Da wir wiederum runde Klammern für Paare verwenden, wurde die Notation im Zitat stillschweigend modernisiert.

Mit Hilfe des geordneten Paares können wir nun eine Multiplikation für Mengen definieren:

Definition (*Kreuzprodukt oder kartesisches Produkt*)

Seien A, B Mengen. Dann ist das *Kreuzprodukt von A und B* , in Zeichen $A \times B$ [A kreuz B], definiert durch:

$$A \times B = \{(a, b) \mid a \in A, b \in B\}.$$

Für $A \times A$ schreiben wir kurz auch A^2 .

Übung

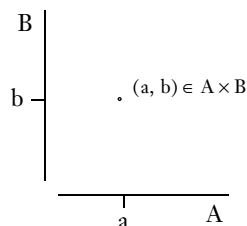
Seien $n, m \in \mathbb{N}$, und seien

$$A = \{0, 1, \dots, n-1\},$$

$$B = \{0, 1, \dots, m-1\}.$$

(In dieser Schreibweise ist $A = \emptyset$ gdw $n = 0$ und analog $B = \emptyset$ gdw $m = 0$.)

Wieviele Elemente hat $A \times B$? Wann ist $A \times B = \emptyset$?



Allgemeiner kann man n -Tupel $(a_1, \dots, a_n)$ aus Objekten sowie $A_1 \times \dots \times A_n$ aus Mengen definieren für $n \in \mathbb{N}$, $n > 2$. Man setzt hierzu:

$$(a_1, \dots, a_n) = (\dots((a_1, a_2), a_3), \dots, a_{n-1}), a_n),$$

$$A_1 \times \dots \times A_n = \{(a_1, \dots, a_n) \mid a_i \in A_i \text{ für alle } 1 \leq i \leq n\}.$$

Weiter sei $A^n = A_1 \times \dots \times A_n$, wobei alle $A_i = A$ sind. Es ist zumeist gefahrlos, A^k mit $A^n \times A^m$ zu identifizieren für alle $n, m, k \in \mathbb{N} - \{0\}$ mit $n + m = k$, obwohl die Elemente von $\mathbb{R}^2 \times \mathbb{R}^2$ etwa von der Form $((x_1, x_2), (x_3, x_4))$, die von $\mathbb{R}^4$ dagegen von der Form (x_1, x_2, x_3, x_4) sind.

Nach Definition gilt für n -Tupel, $n \geq 3$: $(a_1, \dots, a_n) = ((a_1, \dots, a_{n-1}), a_n)$. Solche n -Tupel sind also bestimmte geordnete Paare.

Übung

Zeigen Sie, daß eine Definition des Tripels

$$(a, b, c) = \{\{a\}, \{a, b\}, \{a, b, c\}\}$$

ein Analogon zu (#) nicht erfüllen würde,

d.h. aus $(a, b, c) = (d, e, f)$ folgt hier nicht $a = d, b = e, c = f$.

Dagegen gilt ein Analogon zu (#) für die obige Definition des Tripels als zweifach geschachteltes Paar, $(a, b, c) = ((a, b), c)$, und weiter gilt ein Analogon allgemein für n -Tupel, $n \geq 2$:

$$(\#_n) \quad (a_1, \dots, a_n) = (b_1, \dots, b_n) \quad \text{gdw} \quad a_i = b_i \text{ für alle } 1 \leq i \leq n.$$

Relationen

Als Relationen zwischen Objekten a, b haben wir etwa kennengelernt: $a = b$, $a \in b$, $a \subseteq b$. Die Relationen selbst sind hier $=, \in, \subseteq$. Eine (zweistellige) Relation ist allgemein eine „bestimmte Beziehung“ zwischen zwei Mengen, oder ein „bestimmter Begriff“, der festlegt, wann zwei Objekte in Relation bzgl. dieses Begriffs stehen. Was ist nun genau eine „bestimmte Beziehung“ oder ein „bestimmter Begriff“? Es zeigt sich, daß wir uns hierüber nicht den Kopf zerbrechen müssen; wir würden auch nur Vagheiten aneinanderreihen. Wir definieren einfach: *Eine Relation ist eine Menge von geordneten Paaren*. In dieser Weise wurden Relationen von Charles Peirce, Ernst Schröder und Giuseppe Peano in den beiden letzten Jahrzehnten des 19. Jahrhunderts definiert, wobei dabei der Begriff des geordneten Paares undefiniert blieb.

Diese „Definition mit dem Paukenschlag“ ist das Paradebeispiel für das extensionale Denken der Mengenlehre. Ein Begriff wird mit seinem Umfang identifiziert. Jede Menge von geordneten Paaren R liefert eine Relation, genannt R , die, in wichtigen Fällen wie $\subseteq$, mit einem kontextunabhängigen Namen versehen wird; und zwei Objekte a, b stehen in der Relation R zueinander genau dann, wenn $(a, b) \in R$ gilt.

Die Geschichte des Funktionsbegriffs zeigt, wie schwer sich die Mathematik mit diesem Denken lange Zeit getan hat.

Definition (*Relation*)

Eine Menge R heißt *Relation*, falls jedes $x \in R$ ein geordnetes Paar ist.

Ist A eine Menge und gilt $R \subseteq A \times A$, so heißt R eine *Relation auf A* .

Sind a, b Objekte und gilt $(a, b) \in R$, so schreiben wir hierfür auch $a R b$.

Für jede Relation R kann man eine natürliche Menge A finden, sodaß R eine Relation auf A ist. Hierzu eine Definition.

Definition (*Definitionsbereich und Wertebereich*)

Sei R eine Relation. Wir setzen:

$$\text{dom}(R) = \{ a \mid \text{es existiert ein } b \text{ mit } (a, b) \in R \},$$

$$\text{rng}(R) = \{ b \mid \text{es existiert ein } a \text{ mit } (a, b) \in R \}.$$

$\text{dom}(R)$ heißt der *Definitionsbereich von R* [dom für engl. domain],

$\text{rng}(R)$ heißt der *Wertebereich von R* [rng für engl. range].

Zum Beispiel ist $\text{dom}(A \times B) = A$, $\text{rng}(A \times B) = B$. Jede Relation R ist offenbar eine Relation auf $\text{dom}(R) \cup \text{rng}(R)$.

Die Ordnungsbeziehungen „kleiner“ und „kleinergleich“ auf den Zahlmen-gen $\mathbb{N}$, $\mathbb{Z}$, $\mathbb{Q}$ und $\mathbb{R}$ fassen wir ebenfalls als Relationen auf:

$$<_{\mathbb{N}} = \{ (n, m) \in \mathbb{N} \times \mathbb{N} \mid n \text{ ist kleiner als } m \},$$

$$\leq_{\mathbb{Q}} = \{ (p, q) \in \mathbb{Q} \times \mathbb{Q} \mid p \text{ ist kleiner als } q \text{ oder } p \text{ ist gleich } q \}, \text{ usw.}$$

Eine Relation wie „ $<_{\mathbb{N}}$ “ als ein Objekt zu betrachten, mit dem man weiter operieren kann – etwa läßt sich $\mathcal{P}(<_{\mathbb{N}})$ bilden – ist sicher gewöhnungsbedürftig und zunächst irritierend. Wir haben aber dadurch neben größerer Präzision eine Freiheit der Beschreibung und Manipulation gewonnen, die man schnell lieb gewinnt. Zum Beispiel haben wir folgende Gleichung:

$$\leq_{\mathbb{R}} = <_{\mathbb{R}} \cup \{ (x, x) \mid x \in \mathbb{R} \}.$$

Für die Kleiner-Relationen auf den Zahlen gilt $<_{\mathbb{N}} \subseteq <_{\mathbb{Z}} \subseteq <_{\mathbb{Q}} \subseteq <_{\mathbb{R}}$, und wir unterdrücken deswegen häufig den Index an den Relationen. Beispielsweise gilt dann $(3, 4) \in <$, wobei wir hierfür natürlich zumeist $3 < 4$ schreiben werden. Analoges gilt für $\leq$.

Einige häufig gebrauchte Eigenschaften von Relationen sind:

Definition (*reflexiv, symmetrisch, transitiv*)

Sei R eine Relation auf A .

- (i) R heißt *reflexiv*, falls für alle $x \in A$ gilt $x R x$.
- (ii) R heißt *symmetrisch*, falls für alle $x, y \in A$ gilt:
 $x R y$ *folgt* $y R x$.
- (iii) R heißt *transitiv*, falls für alle $x, y, z \in A$ gilt:
 $x R y$ und $y R z$ *folgt* $x R z$.

Das Trio „reflexiv, symmetrisch, transitiv“ taucht häufiger auf:

Definition (*Äquivalenzrelation*)

Sei R eine Relation auf A . R heißt eine *Äquivalenzrelation auf A* , falls R reflexiv, symmetrisch und transitiv ist.

Für jedes $x \in A$ ist dann die *Äquivalenzklasse* von x bzgl. der Äquivalenzrelation R , in Zeichen x/R [x modulo R], definiert durch:

$$x/R = \{ y \in A \mid x R y \}.$$

Weiter ist A/R [A modulo R] definiert durch $A/R = \{ x/R \mid x \in A \}$.

Ist $y \in x/R$, so heißt y ein *Repräsentant* der Äquivalenzklasse x/R .

$B \subseteq A$ heißt ein *vollständiges Repräsentantensystem für R* , falls für jede Äquivalenzklasse x/R genau ein $y \in B$ existiert mit $y \in x/R$.

Intuitiv bedeutet $x R y$ für eine Äquivalenzrelation R : x und y sind gleich im Sinne von R . Die bevorzugten Zeichen für Äquivalenzrelationen sind deshalb z. B. $\equiv$, $\sim$, $\approx$, ..., die an das Gleichheitssymbol erinnern. Äquivalenzrelationen R auf A sehen die Elemente x von A nur unscharf, wenn nicht $x/R = \{ x \}$ gilt. Die

Identität $\{ (x, x) \mid x \in A \}$ auf A ist der Adler unter den Äquivalenzrelationen auf A , $R = A \times A$ das blinde Huhn.

Ein Beispiel aus der elementaren Zahlentheorie: Für ein festes $m \in \mathbb{N}$, $m \geq 1$, sei die Relation $\equiv_m$ auf $\mathbb{N}$ definiert durch:

$\equiv_m = \{ (a, b) \in \mathbb{N} \times \mathbb{N} \mid a \text{ und } b \text{ haben den gleichen Rest bei Division mit } m \}$.

Dann ist $\equiv_m$ eine Äquivalenzrelation auf $\mathbb{N}$. Für $m = 3$ gilt z. B. $2 \equiv_3 5$, und es gibt genau drei Äquivalenzklassen:

$$0/\equiv_3 = 3/\equiv_3 = 6/\equiv_3 = \dots = \{ 0, 3, 6, 9, \dots \} = \{ k \cdot 3 \mid k \in \mathbb{N} \},$$

$$1/\equiv_3 = 4/\equiv_3 = 7/\equiv_3 = \dots = \{ 1, 4, 7, 10, \dots \} = \{ k \cdot 3 + 1 \mid k \in \mathbb{N} \},$$

$$2/\equiv_3 = 5/\equiv_3 = 8/\equiv_3 = \dots = \{ 2, 5, 8, 11, \dots \} = \{ k \cdot 3 + 2 \mid k \in \mathbb{N} \}.$$

Die Mengen $\{ 0, 1, 2 \}$ und $\{ 0, 1, 5 \}$ sind Beispiele für vollständige Repräsentantensysteme von $\equiv_3$.

Ein Beispiel aus der Analysis: Für $x, y \in \mathbb{R}$ setzen wir

$$x \sim y \text{ falls } |x - y| \in \mathbb{Q}.$$

Dann ist $\sim$ eine Äquivalenzrelation auf $\mathbb{R}$. Die Äquivalenzklassen dieser Relation haben die Form $x/\sim = x + \mathbb{Q} = \{ x + q \mid q \in \mathbb{Q} \}$. In der Analysis sind vollständige Repräsentantensysteme für diese Relation berüchtigt: Sie liefern Beispiele für Teilmengen von $\mathbb{R}$, denen sich keine vernünftige Länge zuordnen läßt (im Fachjargon: sie sind nicht Lebesgue-meßbar, wie Giuseppe Vitali (1875 – 1932) im Jahre 1905 gezeigt hat, vgl. hierzu auch Abschnitt 2, Kapitel 9).

Übung

Sei $\sim$ eine Äquivalenzrelation auf A . Zeigen Sie:

- (i) Für alle $x, y \in A$ gilt: $x \sim y$ gdw $x/\sim = y/\sim$.
- (ii) Für alle $x, y \in A$ gilt: $x/\sim = y/\sim$ oder $x/\sim \cap y/\sim = \emptyset$.
- (iii) Es gilt $\bigcup A/\sim = A$.

Die Äquivalenzklassen bilden also eine Zerlegung von A . Umgekehrt entstehen alle Äquivalenzrelationen aus Zerlegungen von A :

Übung

Sei A eine Menge. Weiter sei $\mathcal{P} \subseteq \mathcal{P}(A)$ eine Zerlegung von A , d. h. es gelte:

- (i) $\emptyset \notin \mathcal{P}$,
- (ii) $\bigcup \mathcal{P} = A$,
- (iii) $P \cap Q = \emptyset$ für alle $P, Q \in \mathcal{P}$ mit $P \neq Q$.

Für $a, b \in A$ sei

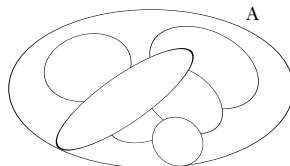
$$a \sim b \text{ falls ein } P \in \mathcal{P} \text{ existiert mit } a, b \in P.$$

Dann gilt:

$\sim$ ist eine Äquivalenzrelation auf A , und es gilt $A/\sim = \mathcal{P}$.

Fassen wir eine Menge A als „Welt“ auf, so ist also eine Äquivalenzrelation auf A nichts anderes als eine Aufteilung der Menge A in bestimmte „Regionen“ oder „Länder“. Und zwei Elemente von A sind äquivalent dann und nur dann, wenn sie in der gleichen Region liegen.

Ist A nichtleer, so hat jede Äquivalenzklasse mindestens ein Element. Andererseits induziert die Zerlegung $\mathcal{P} = \{ \{a\} \mid a \in A \}$ eine Äquivalenzrelation $\sim$, für die jede Äquivalenzklasse einelementig ist: Es gilt $a/\sim = \{a\}$ für alle $a \in A$.



Zerlegung von A in 8 Teile oder
Äquivalenzrelation auf A mit 8 Äquivalenzklassen

Der Funktionsbegriff

Ähnlich wie Cantor „Menge“ definiert hat, könnten wir „Funktion“ definieren als eine Vorschrift, die jedem Element einer bestimmten Menge, dem sogenannten Definitionsbereich der Funktion, eindeutig ein bestimmtes Objekt zuordnet. Der Funktionsbegriff kann aber leicht rein extensional mit Hilfe von geordneten Paaren definiert werden. Funktionen sind danach einfach besondere Relationen, und wir sind nicht mehr auf eine intuitive Beschreibung angewiesen. Unsere Vorstellung und Idee von einer Funktion bleibt aber trotzdem durch die Intuition der „eindeutigen Zuordnung“ bestimmt. Je mehr Konzepte der Leser mit derartigen den Kern der Sache treffenden Vorstellungen versehen und seinem menschlichen Verständnis der Mathematik hinzufügen kann, desto besser. Ohne dieses innere Licht kann man keine Beweise führen und verstehen.

Definition (*Funktion*)

Sei f eine Menge von geordneten Paaren.

f heißt *Funktion*, falls für alle a, b_1, b_2 gilt:

$$(a, b_1) \in f \text{ und } (a, b_2) \in f \text{ folgt } b_1 = b_2. \quad (\text{Rechtseindeutigkeit})$$

Funktionen sind also rechtseindeutige Relationen. Etwas anders formuliert: Eine Relation f ist eine Funktion genau dann, wenn für jedes $a \in \text{dom}(f)$ genau ein b existiert mit $(a, b) \in f$.

Ist f eine Funktion, so schreiben wir wie üblich auch $f(a) = b$ für $(a, b) \in f$.

Definition (*Bild unter einer Funktion*)

Ist f eine Funktion und $f(a) = b$, so heißt b das *Bild von a unter f* .

Die bevorzugten Zeichen für Funktionen sind f, g, h, F, G, H .

Definition (*Funktion von A nach B*)

Seien f eine Funktion und A, B Mengen.

f heißt *Funktion von A nach B*, in Zeichen $f : A \rightarrow B$, falls gilt:

$\text{dom}(f) = A$ und $\text{rng}(f) \subseteq B$.

Diese asymmetrische Behandlung von Definitions- und Wertebereich für „Funktion von ... nach ...“ hat sich als nützlich herausgestellt. Für die Funktion f , die jeder natürlichen Zahl die Null zuordnet, d. h. $f(n) = 0$ für alle $n \in \mathbb{N}$, gilt also $f : \mathbb{N} \rightarrow \mathbb{N}$, obwohl nur ein Wert angenommen wird.

Es hat sehr lange gedauert, bis die einfache abstrakte Definition einer Funktion als Menge von geordneten Paaren formuliert werden konnte. Die bloße Entwicklung des *Begriffs* einer „Funktion“ im 17. Jahrhundert war bereits ein Kraftakt – aus der Antike konnte hier ausnahmsweise einmal nichts übernommen werden. Der Term „Funktion“ selbst geht auf Leibniz zurück. Im 18. Jahrhundert sah man eine Funktion zumeist als eine konkrete Berechnungsvorschrift; für eine Variable wird in eine Funktion eine Zahl eingesetzt, und diese wird in eine andere Zahl umgerechnet. Von Leonhard Euler (1707 – 1783) stammt die Notation „ $f(x)$ “; für ihn sind Funktionen zunächst Kombinationen aus den Grundrechenarten, der Exponentiation und des Logarithmierens, und später allgemeine analytische Ausdrücke, die z. B. auch durch Integration oder implizit durch Gleichungssysteme gegeben sein können; darüber hinaus können Funktionen auch durch unendliche Reihen, Kettenbrüche usw. erzeugt werden. Erst im 19. Jahrhundert setzte sich allmählich der Gedanke einer abstrakten *Korrespondenz* oder *Zuordnung* zwischen Zahlen (Nikolai Lobatschewski 1792 – 1856, Jean Fourier 1768 – 1830, Johann Lejeune Dirichlet 1805 – 1859, Bernhard Riemann 1826 – 1866) und später zwischen beliebigen mathematischen Objekten durch. Ein entscheidender Schritt ist hier Dedekinds Definition, bei der er sich auf Dirichlet beruft, der selber wiederum viel von Fourier übernommen hat:

Dedekind (1887): „Unter einer Abbildung ϕ eines Systems [Menge] S wird ein Gesetz verstanden, nach welchem zu jedem bestimmten Element s von S ein bestimmtes Ding gehört, welches das Bild von s heißt und mit $\phi(s)$ bezeichnet wird ...“

Cantor arbeitet von Beginn an ebenfalls mit einem sehr abstrakten Funktionsbegriff. Die gedankliche, „intern bestimmte“ Möglichkeit, die Elemente zweier Mengen auf bestimmte Art und Weise einander zuzuordnen, wird zur tragenden Säule seiner neuen Konzepte. Für ihn waren Funktionen – wie auch z. B. Ordnungen – keine Mengen, und er hat hier keine Präzisierungsarbeiten übernommen. Von analytischen Ausdrücken und auch „Gesetzen“ ist er aber weit entfernt, und in seinen Schriften lotet er die lichtlosen Tiefen des heutigen abstrakten Funktionskonzepts bereits voll aus, ohne dabei auf eine mengentheoretische Definition von „Funktion“ zurückgreifen zu müssen.

Schließlich werden Funktionen als spezielle Relationen behandelt (insbesondere von Russell, der andererseits ungern Relationen auf geordnete Paare zurückführte). Hausdorff gab dann zum ersten Mal die moderne Definition, bei der Funktionen und geordnete Paare nur Mengen sind und sonst nichts.

Hausdorff (1914): „Zuvor betrachten wir eine Menge P solcher Paare, und zwar von der Beschaffenheit, daß jedes Element a von A in einem und nur einem Paare p von P als erstes Element auftritt. Jedes Element a bestimmt auf diese Weise ein und nur ein Element b , nämlich dasjenige, mit dem es zu einem Paare $p = (a, b)$ verbunden auftritt; dieses durch a bestimmte, von a abhängige, dem a zugeordnete Element bezeichnen wir mit

$$b = f(a)$$

und sagen, daß hiermit in A (d.h. für alle Elemente von A) eine eindeutige Funktion von a definiert sei. Zwei solche Funktionen $f(a)$, $f'(a)$ sehen wir dann und nur dann als gleich an, wenn die zugehörigen Paarmengen P , P' gleich sind, wenn also, für jedes a , $f(a) = f'(a)$ ist.

Umgekehrt ist bei unseren Voraussetzungen ein Element b entweder mit keinem oder einem oder mehreren Elementen a zu einem Paare $p = (a, b)$ verbunden ... “

Verbunden mit unserer Definition von Relation und Funktion sind einige Schreibweisen, die eine Erwähnung verdienen, um möglichen Irritationen vorzubeugen.

Seien $f : A \rightarrow B$, und seien $a \in A$, $b \in B$ mit $f(a) = b$.

Wir setzen $g = f - \{ (a, b) \}$. Dann gilt $g : (A - \{ a \}) \rightarrow B$.

Sei x ein Objekt mit $x \notin A$, und sei y ein beliebiges Objekt.

Sei $h = f \cup \{ (x, y) \}$. Dann gilt $\text{dom}(h) = A \cup \{ x \}$, $\text{rng}(h) = \text{rng}(f) \cup \{ y \}$,

$h : \text{dom}(h) \rightarrow B \cup \{ y \}$, $h(x) = y$.

Die leere Menge $\emptyset$ ist nach Definition eine Funktion mit $\text{dom}(\emptyset) = \emptyset$.

Es gilt $\emptyset : \emptyset \rightarrow B$ für jede Menge B .

Solche Überlegungen gehören zu den Tonleiterübungen der Mengenlehre, und sollten auch als solche aufgefaßt werden.

Einfache Eigenschaften einer Funktion

Die folgenden Eigenschaften *injektiv*, *surjektiv*, *bijektiv* einer Funktion sind für die Mengenlehre von zentraler Bedeutung.

Definition (*injektiv*, *surjektiv*, *bijektiv*)

Sei f eine Funktion.

- (i) f heißt *injektiv*, falls für alle a_1, a_2, b gilt:
 $f(a_1) = b$ und $f(a_2) = b$ folgt $a_1 = a_2$. (*Linkseindeutigkeit*)
- (ii) f heißt *surjektiv nach B* , falls gilt $\text{rng}(f) = B$, d.h.
für alle $b \in B$ existiert ein $a \in \text{dom}(f)$ mit $f(a) = b$.
 f heißt *surjektiv von A nach B* , falls $f : A \rightarrow B$ und f surjektiv nach B .
- (iii) f heißt *bijektiv von A nach B* , falls gilt:
 f ist injektiv und surjektiv von A nach B .

Man beachte, daß „injektiv“ keine Erwähnung von Definitions- und Wertebereich verlangt.

Wir schreiben auch „ $f : A \rightarrow B$ surjektiv“ für „ f ist surjektiv von A nach B “. Analog für bijektiv.

Die drei Konzepte kann man sich leicht vor Augen führen:

Anschaulich bedeutet

$f : A \rightarrow B$	injektiv:	kein Wert (in B) wird mehrfach angenommen,
	surjektiv:	jeder Wert in B wird angenommen,
	bijektiv:	vollständige Paarbildung zwischen den Elementen von A und B .

Nur Bijektionen behandeln ihren Definitions- und Wertebereich vollkommen symmetrisch.

Fraenkel (1959) über Bijektionen: „Grundlegend für alles Folgende ist der Begriff der *Abbildung* [bei Fraenkel hier = Bijektion]. Wird jedem Element m einer Menge M ein einziges Element n einer Menge N zugeordnet, so spricht man von einer eindeutigen Zuordnung von Elementen aus N zu den Elementen von M ; dabei kann natürlich verschiedenen m das nämliche n entsprechen, wie es z. B. der Fall ist, falls n den Vater von m bedeutet. Ist aber die Zuordnung überdies auch eindeutig in der umgekehrten Richtung, d. h. entspricht auch jedem n ein einziges m , wie es z. B. für die Zuordnung zwischen Ehemännern und Ehefrauen in einer monogamen Gesellschaftsordnung zutrifft, so heißt die Zuordnung *eindeutig*. Sie wird dann auch eine *Abbildung zwischen M und N genannt* [= Bijektion zwischen M und N], und die hiermit ausgedrückte Symmetrie zwischen beiden Mengen hinsichtlich der Zuordnung ist offenbar begründet.“

Satz

- (i) Ist $f : A \rightarrow B$, so ist $f : A \rightarrow \text{rng}(f)$ surjektiv.
- (ii) Ist $f : A \rightarrow B$ injektiv, so ist $f : A \rightarrow \text{rng}(f)$ bijektiv.
- (iii) Ist $f : A \rightarrow B$ bijektiv, so existiert ein $g : B \rightarrow A$ bijektiv.

Beweis

zu (i) und (ii): Die Behauptungen folgen unmittelbar aus den Definitionen.

zu (iii): Wir setzen

$$g = \{ (b, a) \mid (a, b) \in f \}.$$

– Dann ist $g : B \rightarrow A$ bijektiv.

Um den Leser weiter mit dem Funktionsbegriff und seinen ungewöhnlichen Notationen vertraut zu machen, halten wir noch einige Merkmale fest.

Nach Definition einer Funktion gilt: Ist f eine Funktion, und ist $g \subseteq f$, so ist auch g eine Funktion. Ist f injektiv, so vererbt sich die Injektivität auf jedes $g \subseteq f$.

Sind f und g zwei Funktionen, und gilt $f(x) = g(x)$ für alle $x \in \text{dom}(f) \cap \text{dom}(g)$, so ist auch $f \cup g$ eine Funktion. Ist F eine Menge von Funktionen, und gilt $f(x) = g(x)$ für alle $f, g \in F$ und alle $x \in \text{dom}(f) \cap \text{dom}(g)$, so ist auch $\bigcup F$ eine

Funktion. Verschiedene Funktionen lassen sich also zusammenbauen, wenn sie sich auf ihren Definitionsbereichen verträglich verhalten. Weiter ist immer auch $\bigcap F$ für $F \neq \emptyset$ eine Funktion, denn es gilt $\bigcap F \subseteq f$ für alle $f \in F$.

Neben diesen definitorischen Spielereien gibt es solche für die Anschauung: Ist $f: A \rightarrow B$ eine Injektion, so kann man sich f als eine Operation vorstellen, die den Elementen von A Namen aus der Menge B zuweist: Jedes $x \in A$ erhält den Namen – oder den Zettel – $f(x) \in B$. Injektionen lassen sich so als Umbenennungen auffassen, wenn man vereinbart, daß der ursprüngliche Name von x nichts anderes ist als x selbst. Hintereinandergeschaltete Injektionen entsprechen dann wiederholten Umbenennungen. Eine nicht injektive Funktion $f: A \rightarrow B$ dagegen kann man sich als eine Operation vorstellen, die bestimmte Elemente von A miteinander identifiziert oder verklebt. Ist $f: A \rightarrow B$ eine Funktion, so ist $\sim$ eine Äquivalenzrelation auf A , wobei $x \sim y$, falls $f(x) = f(y)$.

Einfache Operationen mit Funktionen

Die einfachste Operation, die aus einer Funktion eine weitere Funktion macht, ist die Einschränkung der Funktion auf einen Teil des Definitionsbereichs:

Definition (*Einschränkung einer Funktion*)

Sei $f: A \rightarrow B$ und $C \subseteq A$. Dann setzen wir

$$f|C = \{ (a, b) \in f \mid a \in C \}.$$

$f|C$ heißt die *Einschränkung von f auf C* .

Sind f, g Funktionen mit $f \subseteq g$, so heißt g eine *Fortsetzung von f* .

Ist $f: A \rightarrow B$ und $C \subseteq A$, so gilt $f|C: C \rightarrow B$, und f ist eine Fortsetzung von $f|C$. Eine Funktion g ist Fortsetzung einer Funktion f genau dann, wenn $g|_{\text{dom}(f)} = f$.

Eine weitere Operation ist uns im Beweis oben schon begegnet, nämlich die „Umkehrung“ einer Funktion. Hier brauchen wir die Injektivität (Linkseindeutigkeit) der Ausgangsfunktion, damit die Rechtseindeutigkeit der Umkehrung gewährleistet ist.

Definition (*Umkehrfunktion*)

Sei f eine injektive Funktion. Dann ist die *Umkehrfunktion von f* , in Zeichen f^{-1} , definiert durch

$$f^{-1} = \{ (b, a) \mid (a, b) \in f \}.$$

Offenbar ist f^{-1} wieder eine injektive Funktion, und es gilt $\text{dom}(f^{-1}) = \text{rng}(f)$, $\text{rng}(f^{-1}) = \text{dom}(f)$. Weiter gilt $(f^{-1})^{-1} = f$.

Definition (*Verkettung oder Verknüpfung von Funktionen*)

Seien f, g Funktionen. Dann ist die *Verkettung* $g \circ f$ der Funktionen f und g [gelesen: g nach f] definiert durch:

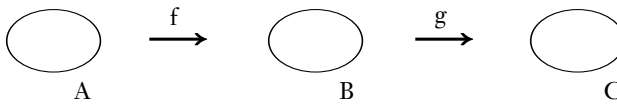
$$g \circ f = \{ (a, c) \mid a \in \text{dom}(f) \text{ und } (f(a), c) \in g \}.$$

Übung

Seien f, g Funktionen. Dann gilt:

- (i) $g \circ f$ ist eine Funktion, $\text{dom}(g \circ f) \subseteq \text{dom}(f)$, $\text{rng}(g \circ f) \subseteq \text{rng}(g)$.
- (ii) Ist $f : A \rightarrow B$ und $g : B \rightarrow C$, so ist $g \circ f : A \rightarrow C$.

Seien $f : A \rightarrow B$, $g : B \rightarrow C$, und sei $h = g \circ f : A \rightarrow C$.



Nach Definition haben wir $h(a) = g(f(a))$ für alle $a \in A$. Die Verkettung h von f und g ist also die Hintereinanderausführung von f und g : Wir starten bei a , bilden $b = f(a)$ und setzen dann $h(a) = g(b) = g(f(a))$.

Weiter gilt in dieser Situation stets $g \circ f = g \mid \text{rng}(f) \circ f$.

Übung

Seien $f : A \rightarrow B$ und $g : B \rightarrow C$ Funktionen und sei $h = g \circ f$.

- (i) Sind f, g injektiv, so ist auch h injektiv.
- (ii) Ist $g \mid \text{rng}(f)$ surjektiv nach C , so ist h surjektiv nach C .
- (iii) Sind $f : A \rightarrow B$ und $g : B \rightarrow C$ bijektiv, so ist $h : A \rightarrow C$ bijektiv.

Die Verkettung ist eine assoziative Operation, was man sich anhand eines Diagramms wie oben anschaulich machen kann:

Übung (*Assoziativgesetz für die Verknüpfung von Funktionen*)

Seien f, g, h Funktionen.

Dann gilt $(h \circ g) \circ f = h \circ (g \circ f)$.

Wir können also Klammern weglassen und $h \circ g \circ f$ schreiben.

Für die einfachsten Funktionen brauchen wir noch eine Bezeichnung:

Definition (*Identität*)

Sei A eine Menge. Dann ist die *Identität auf A* , in Zeichen id_A , definiert durch:

$$\text{id}_A = \{ (a, a) \mid a \in A \}.$$

Es gilt $\text{id}_A : A \rightarrow A$. Offenbar ist id_A eine Bijektion von A nach A . Nach Definition ist $\text{id}_\emptyset = \emptyset$.

Übung

- (a) Sei $f : A \rightarrow B$ bijektiv.
Dann gilt $f^{-1} \circ f = \text{id}_A$ und $f \circ f^{-1} = \text{id}_B$.
- (b) Seien f, g injektiv. Dann gilt $(g \circ f)^{-1} = f^{-1} \circ g^{-1}$.

Der folgende Satz bringt ein nützliches Kriterium für die Injektivität oder Surjektivität einer Funktion.

Satz (*Verkettung zur Identität*)

Sei $f : A \rightarrow B$ eine Funktion.

- (a) Ist $A \neq \emptyset$, so sind die folgenden Aussagen äquivalent:
- (i) $f : A \rightarrow B$ ist injektiv.
 - (ii) Es existiert eine Funktion $g : B \rightarrow A$ mit $g \circ f = \text{id}_A$.
- (b) Die folgenden Aussagen sind äquivalent:
- (i) $f : A \rightarrow B$ ist surjektiv.
 - (ii) Es existiert eine Funktion $g : B \rightarrow A$ mit $f \circ g = \text{id}_B$.

Beweis

zu (a):

(i) $\hookrightarrow$ (ii):

Seien $B' = \text{rng}(f)$ und $a \in A$ beliebig. Definiere $g : B \rightarrow A$ durch
 $g|_{B'} = f^{-1}$, $g(b) = a$ für $b \notin B'$.
 Dann ist $g \circ f = \text{id}_A$.

(ii) $\hookrightarrow$ (i):

Sei $f(a_1) = b = f(a_2)$ für $a_1, a_2 \in A$. Dann ist $g(b) = g(f(a_1)) = \text{id}_A(a_1) = a_1$ und $g(b) = g(f(a_2)) = \text{id}_A(a_2) = a_2$, also $a_1 = a_2$. Also ist f injektiv.

zu (b):

(i) $\hookrightarrow$ (ii):

Wir definieren $g : B \rightarrow A$ durch
 $g(b) = \text{„ein } a \in A \text{ mit } f(a) = b\text{“}$ für alle $b \in B$.
 [Ein $a \in A$ mit $f(a) = b$ existiert wegen $\text{rng}(f) = B$ nach Voraussetzung.]
 Dann gilt $g : B \rightarrow A$ und nach Definition von g ist $f(g(b)) = b$ für alle $b \in B$. Also $f \circ g = \text{id}_B$.

(ii) $\hookrightarrow$ (i):

Sei $b \in B$. Wegen $f \circ g = \text{id}_B$ ist $f(g(b)) = b$. Also ist $b \in \text{rng}(f)$,
 — denn $f(a) = b$ für $a = g(b)$. Also ist $f : A \rightarrow B$ surjektiv.

Dem Leser ist hier vielleicht die folgende Zeile aufgefallen:

$g(b) = \text{„ein } a \in A \text{ mit } f(a) = b\text{“}$.

Manche Leser werden dieses „ein“ vielleicht seltsam finden, andere werden fragen, was daran besonderes sein soll. Die einzige Besonderheit ist, daß hier ein

sehr abstrakter Prozeß am Werk ist (und manche Leser werden sich mittlerweile schon an unbegrenzte Abstraktion gewöhnt haben): Wir wissen, daß für jedes b in B mindestens ein a in A mit $f(a) = b$ existiert. Für die Definition von g müssen wir aber derartige Zeugen a für alle $b \in B$ herauspicken. Überspitzt gefragt: Wer garantiert uns, daß es einen auf ganz B operierenden „Zufallsgenerator“ gibt, der uns diese Auswahlarbeit „liefere, gegeben b , irgendein für dieses b geeignetes a “ abnimmt? Die Antwort aus heutiger Sicht ist: Innerhalb eines genügend reichhaltigen naiven Mengenbegriffs liegt kein Problem vor, es gibt immer solche a 's, also nehmen wir einfach welche und definieren g . Innerhalb der axiomatischen Behandlung der Mengenlehre zeigt sich, daß man das Funktionieren solcher platonisch zweifellos gerechtfertigter Auswahlprozesse durch ein eigenes und sehr starkes Axiom absichern muß.

Im folgenden deuten wir ähnliche Situationen durch verwandte Schreibweisen wie „ein ...“ an, und betrachten diese sehr abstrakte Form der Mengenbildung als interessant aber unproblematisch. Begegnet ist sie uns an versteckter Stelle schon zuvor: Will man beweisen, daß jede Äquivalenzrelation ein vollständiges Repräsentantensystem besitzt, ist ein ähnliches „ein ...“ nötig.

Bilder und Urbilder

Häufig gebraucht wird das Aufsammeln von Bildern und Urbildern:

Definition (*Bild und Urbild*)

Sei f eine Funktion, und seien A, B Mengen. Wir setzen:

$$f''A = \{ f(a) \mid a \in \text{dom}(f) \cap A \},$$

$$f^{-1}B = \{ a \in \text{dom}(f) \mid f(a) \in B \}.$$

$f''A$ heißt *das Bild von A unter f* [gelesen: f zwei Strich von A].

$f^{-1}B$ heißt *das Urbild von B unter f*.

Man beachte, daß die Urbildoperation f^{-1} auch für nicht injektive Funktionen definiert ist.

Gebräuchlich sind auch die Bezeichnungen $f[A]$ und $f^{-1}[B]$ für $f''A$ und $f^{-1}B$. Es gilt $f''A \subseteq \text{rng}(f)$ und $f^{-1}B \subseteq \text{dom}(f)$ für alle Mengen A, B .

Übung

Sei $f : A \rightarrow B$ eine Funktion. Dann gilt:

- (i) $f''f^{-1}X \subseteq X$ für alle $X \subseteq B$,
- (ii) $f''f^{-1}X = X$ für alle $X \subseteq \text{rng}(f)$,
- (iii) $f^{-1}f''X \supseteq X$ für alle $X \subseteq A$,
- (iv) $f^{-1}f''X = X$ für alle $X \subseteq A$, falls f injektiv ist.

Übung

Sei $f : A \rightarrow B$ eine Funktion. Dann gilt:

- (i) $f''(X - Y) \supseteq f''X - f''Y$ für alle $X, Y \subseteq A$,
- (ii) $f''(X - Y) = f''X - f''Y$ für alle $X, Y \subseteq A$, falls f injektiv ist,
- (iii) $f^{-1}''(X - Y) = f^{-1}''X - f^{-1}''Y$ für alle $X, Y \subseteq B$,
- (iv) $f'' \cap \mathcal{X} \subseteq \bigcap \{f''X \mid X \in \mathcal{X}\}$ für alle $\mathcal{X} \subseteq \mathcal{P}(A)$,
- (v) $f'' \cap \mathcal{X} = \bigcap \{f''X \mid X \in \mathcal{X}\}$ für alle $\mathcal{X} \subseteq \mathcal{P}(A)$, falls f injektiv ist,
- (vi) $f'' \cup \mathcal{X} = \bigcup \{f''X \mid X \in \mathcal{X}\}$ für alle $\mathcal{X} \subseteq \mathcal{P}(A)$,
- (vii) $f^{-1}'' \cap \mathcal{X} = \bigcap \{f^{-1}''X \mid X \in \mathcal{X}\}$ für alle $\mathcal{X} \subseteq \mathcal{P}(B)$,
- (viii) $f^{-1}'' \cup \mathcal{X} = \bigcup \{f^{-1}''X \mid X \in \mathcal{X}\}$ für alle $\mathcal{X} \subseteq \mathcal{P}(B)$.

Für die Urbildoperation gilt also stets Gleichheit. Woran liegt das? Zunächst kann man sich vor Augen führen, daß Injektivität von f eine gute Eigenschaft ist, da dann $f(x)$ einfach als ein „neuer Name“ für x angesehen werden kann, während eine derartige Umbenennung bei nichtinjektiven Funktionen zu Identitätskrisen führen würde. Damit ist dann die Gleichheit in (ii) und (v) klar. Faßt man nun f^{-1} als eine „mehrdeutige Funktion“ auf $\text{rng}(f)$ auf, die einem x u. U. mehrere Werte zuweist, nämlich alle y mit $f(y) = x$, so gilt „Injektivität“ für diese mehrdeutige Funktion: Die Mengen der Werte, die zu verschiedenen x_1, x_2 gehören, sind disjunkt (denn f ist als Funktion rechtseindeutig). f^{-1} kann man also als eine Umbenennung auffassen, bei der Objekte u. U. mehrere neue Namen erhalten, jedoch dieselben Namen niemals an verschiedene Objekte vergeben werden.

Der Leser, der immer noch nicht erschöpft ist, kann versuchen, für Relationen R, S die Operationen

$$\begin{aligned} R^{-1} &= \{ (b, a) \mid (a, b) \in R \}, \\ R \circ S &= \{ (a, c) \mid \text{es gibt ein } b \text{ mit } (a, b) \in R \text{ und } (b, c) \in S \}, \\ R''A &= \{ b \mid (a, b) \in R \text{ für ein } a \in A \}, \\ R^{-1}''B &= (R^{-1})''B = \{ a \mid (a, b) \in R \text{ für ein } b \in B \} \end{aligned}$$

in ähnlicher Weise zu untersuchen, wie wir es für den Funktionsbegriff getan haben. (Diese Definitionen rechtfertigen auch die Schreibweise $f^{-1}''B$ für nicht-injektive Funktionen f ; die Relation f^{-1} ist nun immer definiert, da eine Funktion f eine Relation ist. Obige Verknüpfung $\circ = \circ_{\text{Rel}}$ von Relationen ist invers zur Verknüpfung $\circ = \circ_{\text{Fun}}$ von Funktionen, d. h. $g \circ_{\text{Fun}} f = f \circ_{\text{Rel}} g$, was aber i. a. problemlos ist, sodaß der Index an $\circ$ entfallen kann.)

Nach dieser zuletzt recht trockenen Anreicherung unserer Sprache im Umfeld des Funktionsbegriffs nähern wir uns im folgenden Kapitel der zentralen Idee der Mengenlehre über eine spielerische Frage ...

P. R. Halmos über den mengentheoretischen Funktionsbegriff

„Dementsprechend werden die Bezeichnungen *Abbildung*, *Transformation*, *Zuordnung* sowie *Operator* (und viele andere) zuweilen als Synonyma für Funktion verwendet ...

Die oben aufgezählten Synonyma für ‚Funktion‘ suggerieren alle irgendeine Tätigkeit als Nebenbedeutung. Deshalb sind manche Mathematiker unzufrieden mit unserer Definition, derzufolge eine Funktion nichts (im dynamischen Sinne) tut, sondern einfach (im statischen Sinne) vorhanden ist. Diese Unzufriedenheit äußert sich in einem abweichenden Gebrauch des Vokabulars: Funktion wird für das undefinierte Objekt reserviert, das irgendwie aktiv ist, und die Menge geordneter Paare, die wir als Funktion bezeichnet haben, wird dann der *Graph* der Funktion genannt.“

(P. R. Halmos 1972, „*Naive Mengenlehre*“)

4. Größenvergleiche

Wie vergleicht man zwei große Haufen Nüsse?

Gemeint ist: Wie stellt man fest, welcher der beiden Nußhaufen mehr Nüsse enthält?

Der Praktiker wird sagen: Wiegen! Da wir nach einer möglichst elementaren Lösung suchen, denken wir uns, wir befänden uns in einer prähistorischen Zeit, wo man keine Waagen kennt, und auch keine allzu großen natürlichen Zahlen.

Also ist auch die Idee, die beiden Haufen H_1 und H_2 abzuzählen und so die Anzahlen der Nüsse festzustellen, ungeeignet.

Einem einigermaßen begabten prähistorischen Sammler könnte aber folgender Algorithmus einfallen:

Der Algorithmus des Abtragens

- (1) Nimm je eine Nuß von H_1 und H_2 weg und lege sie beiseite.
- (2) Wiederhole diese Paarbildung, bis einer der beiden Haufen aufgebraucht ist.

Die Resultate:

- (I) H_1 ist am Ende aufgebraucht.
- (II) H_2 ist am Ende aufgebraucht.
- (III) H_1 und H_2 sind am Ende aufgebraucht.

entsprechen genau den Fällen:

- (I) H_1 hat weniger oder ebensoviele Nüsse wie H_2 .
- (II) H_2 hat weniger oder ebensoviele Nüsse wie H_1 .
- (III) H_1 und H_2 haben gleichviele Nüsse.

Diese Entsprechung bezeichnen wir als die *Korrektheit des Algorithmus*, von der wir intuitiv überzeugt sind.

Der Algorithmus erzeugt eine Funktion: Sei f die Menge aller Paare $(a, b) \in H_1 \times H_2$, die während des Abtragens entstehen.

Im Fall (I) ist $f : H_1 \rightarrow H_2$ injektiv.

Im Fall (II) sei $g = f^{-1} = \{ (b, a) \mid (a, b) \in f \}$. Dann ist $g : H_2 \rightarrow H_1$ injektiv.

Im Fall (III) ist $f : H_1 \rightarrow H_2$ bijektiv.

Aus der Korrektheit des Algorithmus folgt, daß das Ergebnis nicht vom Verlauf der Paarbildung abhängt, d.h. jede mögliche Durchführung – jede mögliche entstehende Funktion f – liefert das gleiche Resultat.

Wir setzen:

- $|H_1| \leq |H_2|$, falls (I) eintritt,
- $|H_2| \leq |H_1|$, falls (II) eintritt,
- $|H_1| = |H_2|$, falls (III) eintritt.

[gelesen: Betrag H_1 kleinergleich Betrag H_2 , usw.]

Wir betonen noch einmal, daß wir hier keine natürlichen Zahlen verwenden – die Begriffe „mehr“, „weniger“ und „gleichviel“ beruhen damit nicht auf dem Zahlbegriff.

Zusammenfassend können wir festhalten:

Größenvergleich und Existenz von Funktionen	
(1)	$ H_1 \leq H_2 $ gdw „es existiert ein injektives $f: H_1 \rightarrow H_2$ “.
(2)	$ H_2 \leq H_1 $ gdw „es existiert ein injektives $f: H_2 \rightarrow H_1$ “.
(3)	$ H_1 = H_2 $ gdw „es existiert ein bijektives $f: H_2 \rightarrow H_1$ “.

Es gilt weiter:

- (3') $|H_1| = |H_2|$ gdw „es existieren $f: H_1 \rightarrow H_2$ injektiv und $g: H_2 \rightarrow H_1$ injektiv“.

Die Richtung von links nach rechts ist klar nach (3).

Für die andere Richtung argumentieren wir so:

Aus der Existenz von f folgt nach (1), daß $|H_1| \leq |H_2|$.

Aus der Existenz von g folgt nach (2), daß $|H_2| \leq |H_1|$.

Also gilt (I) und (II), und damit (III), also $|H_1| = |H_2|$.

Wir werden obige Tabelle zur *Definition* von $|M| \leq |N|$ und $|M| = |N|$ verwenden für beliebige Mengen M und N . (3') – eine Folgerung unserer Korrektheits-Annahme – müssen wir dann beweisen, und dies wird unsere erste nichttriviale Aufgabe sein: Der Satz von Cantor-Bernstein.

Der Leser wird vielleicht fragen: Wie, wenn wir mit $<$ arbeiten statt mit $\leq$? Dann erhalten wir z.B., daß H_1 echt weniger Nüsse hat als H_2 , wenn es ein Abtragen gibt, das H_1 leert, aber bei H_2 einen Rest hinterläßt. Wir könnten also obige Tabelle ergänzen:

- (+) $|H_2| < |H_1|$ gdw „es existiert ein injektives, nicht surjektives $f: H_2 \rightarrow H_1$ “.

Das ist für den Abtragealgorithmus für Nußhaufen sicher richtig, aber diese Zeile könnte man nicht sinnvoll zu einer Definition von $<$ für beliebige Mengen M, N verwenden: Unter dieser Definition würde nicht gelten, daß $|M| < |N|$ dasselbe ist wie „ $|M| \leq |N|$ und $|M| \neq |N|$ “. Der Leser kann sich das jetzt schon am Beispiel von $M = N = \mathbb{N}$ klarmachen. Wir kommen im Kapitel über Unendlichkeit ausführlich darauf zurück. Es erweist sich als eine Besonderheit der Nußhaufen und der endlichen Mengen, daß (+) das zu $\leq$ und $=$ gehörige $<$ definiert.

Der Rückgriff auf Nahrungsmittel zur Motivation des Größenvergleichs hat eine lange Tradition. So schreibt Hessenberg 1906 in seinem Buch „Grundbegriffe der Mengenlehre“:

Hessenberg (1906): „Denken wir uns als ganz konkretes Beispiel etwa einen Korb Äpfel und einen Korb Birnen. Um zu prüfen, ob beide gleichviel Stücke enthalten, können wir so verfahren: Wir ergreifen mit der linken Hand einen Apfel, mit der rechten eine Birne und legen jedes Stück aus seinem Korb heraus. Dieses Verfahren wiederholen wir, solange es geht. Es wird nach einer endlichen Anzahl von Wiederholungen zu einem Ende führen, und zwar entweder dadurch, daß keine Äpfel mehr da sind, oder dadurch, daß keine Birnen mehr vorhanden sind, oder drittens dadurch, daß weder Birnen noch Äpfel mehr übrig bleiben. Im ersten Fall ist die Anzahl der Birnen, im zweiten die der Äpfel die größere, im dritten Fall sind die Anzahlen einander gleich.“

Hausdorff denkt das Beispiel von Hessenberg kulturgeschichtlich weiter, und gelangt über die Einführung von vermittelnden Objekten schließlich zu Zeichensystemen und abstrakten Zahlen:

Hausdorff (1914): „Wenn man eine Menge von Äpfeln mit einer Menge von Birnen auf die Anzahl der Gegenstände vergleichen will, so geschieht dies auf dem primitiven Standpunkt in der Weise, daß man einen Apfel mit einer Birne zusammenlegt, und dieses Verfahren bis zu seinem Ende fortsetzt...

Wenn aber die Äpfel und Birnen sich an verschiedenen Orten befinden und ein Transport der einen Menge zu der anderen mit Schwierigkeiten verbunden ist, so wird der erfinderische Menscheng Geist auf der nächsthöheren Stufe sich einer vermittelnden Menge bequem transportabler Gegenstände, seien es Steine, Muscheln, Holzstücke, bedienen und [durch Vergleich von Äpfeln mit den Gegenständen und Vergleich der Gegenstände mit den Birnen die Anzahlen der Äpfel und Birnen vergleichen]. Endlich wird aber auch dieser Erdenrest noch überwunden und an Stelle der vermittelnden Menge tritt ein System gesprochener, geschriebener oder gedachter Zeichen, der Zahlzeichen 1, 2, Das Vergleichen wird damit zum Zählen, und äquivalente [gleichgroße] Mengen erhalten nun eine gemeinsame Eigenschaft, die Anzahl ihrer Elemente.

Diese Bemerkungen, die weder nach der psychologischen, noch nach der kulturgeschichtlichen Seite hin irgendwelchen Anspruch erheben wollen, sollen nur verständlich machen, daß die Äquivalenz [Möglichkeit einer vollständigen, symmetrischen Paarbildung] die natürliche Grundlage der Vergleichen von Mengen ist und daß mit ihrer Hilfe sogar der anscheinend paradoxe Versuch unternommen werden konnte, auch unendliche Mengen zu zählen.“

Die Idee des Vergleichs durch Paarbildung reicht sehr weit in den Brunnen der Vergangenheit zurück. Wir geben zuerst die offiziellen Definitionen, die am Ende eines langen Schlafes und eines ebenso ausgedehnten Erwachens stehen, und die dort dann einen Meilenstein in der Geschichte der Mathematik markieren. Anschließend zeichnen wir nicht gerade die Höhlenmalereien, aber doch einige frühe Skizzen der Paarbildung nach, die uns die Geschichte überliefert hat.

Größenvergleich zweier Mengen: „gleichgroß“

Motiviert durch den Algorithmus des Abtragens definieren wir nun einen rein mathematischen Begriff. Dieser Begriff der Gleichmächtigkeit zweier Mengen erweist sich als derart tragfähig und ebenmäßig, daß sich auf seiner Grundfläche eine Pyramide errichten läßt, die einen Vergleich mit einem altägyptischen Original weder nach Konstruktion noch nach Wirkung zu scheuen braucht.

Definition ($|A| = |B|$; erste Fundamentaldefinition der Mengenlehre)

Seien A und B Mengen. Dann ist A gleichmächtig zu B , in Zeichen $|A| = |B|$, falls gilt:

Es existiert ein bijektives $f : A \rightarrow B$.

Cantor (1878): „Wenn zwei wohldefinierte Mannigfaltigkeiten M und N sich eindeutig und vollständig, Element für Element, einander zuordnen lassen (was, wenn es auf eine Art möglich ist, immer auch noch auf viele andere Weisen geschehen kann), so möge für das folgende die Ausdrucksweise gestattet sein, daß diese Mannigfaltigkeiten gleiche Mächtigkeit haben, oder auch, daß sie äquivalent sind.“

Der zweite Fundamentalbegriff der Mengenlehre ist der der Wohlordnung, ebenfalls von Cantor aufgestellt und in den Mittelpunkt gerückt. Wir werden ihn im zweiten Abschnitt definieren und unter die Lupe nehmen. Er ist die Senkrechte zur Grundfläche der Gleichmächtigkeit und erlaubt uns sichere Konstruktionen in schwindelerregende Höhen.

Übung

Seien A, B, C Mengen. Dann gilt:

- (i) $|A| = |A|$, (Reflexivität)
- (ii) $|A| = |B|$ folgt $|B| = |A|$, (Symmetrie)
- (iii) $|A| = |B|$ und $|B| = |C|$ folgt $|A| = |C|$. (Transitivität)

Fraenkel (1959): „Wir setzen nun fest:

Definition 2. Existiert eine Abbildung [= Bijektion] der Menge N auf die Menge M , so heißt M äquivalent der Menge N , in Zeichen $M \sim N$ [$|M| = |N|$].

Offenbar ist jede Menge sich selbst äquivalent ($M \sim M$), aus $M \sim N$ folgt $N \sim M$, und aus den Beziehungen $M \sim N$, $N \sim P$ folgt $M \sim P$. Kurz, die Äquivalenz ist eine reflexive, symmetrische und transitive Relation. Im allgemeinen – nämlich wenn M mehr als ein Element enthält – gibt es verschiedene Abbildungen zwischen zwei äquivalenten Mengen ...“

Übung

Seien M, A, B Mengen, $A \cap B = \emptyset$ und es gelte $|M| = |M \cup A|$, $|M| = |M \cup B|$. Dann gilt $|M| = |M \cup A \cup B|$.

Zur Geschichte des Vergleichs durch Paarbildung

Welche Sammler der grauen Vorzeit Nüsse, Schafe, Untergebene, Skalps, Muscheln oder ähnliches durch Paarbildung miteinander verglichen haben, bleibt uns unerforschlich. Aber es gibt dann bereits im Mittelalter recht abstrakte Paarbildungen, entstanden nicht im Kontext der irdischen Zwänge, sondern beim Ringen um den Begriff der Unendlichkeit, beim scholastischen Argumentieren über die Existenz und Nichtexistenz unendlicher Objekte und Ideen in Raum, Zeit, Mathematik und Gott. (Wir behandeln die Begriffe endlich und unendlich mathematisch im übernächsten Kapitel, hier genügt ein naives Verständnis der Begriffe vollauf.)

Die folgende Darstellung stützt sich zuvörderst auf die Arbeiten von Helmuth Gericke, insbesondere [Gericke 1977], aus der im folgenden zitiert wird. (Siehe weiter auch [Maier 1949].)

Die griechische Philosophie erkundete seit Anaximander (~ 611 – 546 v. Chr.) die Idee des Unendlichen (des *ἄπειρον*, d. h. wörtlich: des Unbegrenzten), und Aristoteles (384 – 322 v. Chr.) behandelt das Unendliche ausführlich in seiner *Physik*. Die Scholastik des Mittelalters kommentiert Aristoteles hauptberuflich, nach und nach werden aber in den Kommentaren des 13. und 14. Jahrhunderts eigene Gedanken entwickelt, es beginnt eine Emanzipierung durch iterierte Kommentierung, und die Position des konservierenden Papageis wird schließlich als unbefriedigend empfunden und abgelegt.

Seit der Antike herrschte der heute noch nachvollziehbare Zweifel daran, daß ein Kontinuum aus einzelnen Punkten zusammengesetzt sein könne; des weiteren herrschte die heute schwer zu verstehende Verwirrung, daß ein Unendliches mit einem anderen Unendlichen völlig identisch sein müsse. Roger Bacon (~ 1210 – 1292) geht der Frage der möglichen Punktartigkeit des Kontinuums nach, und bildet dabei die untere Seite eines Quadrats auf die Diagonale des Quadrats bijektiv ab, indem er jeden Punkt der Seite mit dem darüberliegenden Punkt der Diagonalen paart. Damit wäre, so Bacon, die Grundseite identisch mit der Diagonalen, ein Widerspruch – also kann ein Kontinuum nicht aus Punkten zusammengesetzt sein.

Hier besonders interessant und kulturgeschichtlich sehr einflußreich ist das Viergestirn Wilhelm von Ockham (~ 1280/1285 – 1347/1349), sein Schüler Johannes Buridan (~ 1295 – 1358), und dessen Schüler Nikolaus von Oresme (~ 1320 – 1382) und Albert von Sachsen (~ 1316/1325 – 1390). Innerhalb dieser Gruppe entstehen neue Aristoteleskommentare, und schließlich finden sich bei ihrem jüngsten Mitglied schockierend moderne Einblicke. Euklid gibt das entscheidende Stichwort mit seinem geometrischen Slogan „was sich deckt, ist gleich“. Albert von Sachsen greift den Gedanken auf und definiert die „Deckungsgleichheit“ für *multitudines*, also für Mengen: „Wenn zwei Mengen sich so verhalten, daß jeder Einheit der einen eine Einheit der anderen entspricht, dann ist die eine weder größer noch kleiner als die andere. Das erscheint als an sich gesichert, da ja die eine die andere nicht übersteigt.“ Mengenlehre im 14. Jahrhundert! Albert gibt ein Beispiel: Gegeben sei ein zu einer Seite unendlich langer Balken mit einem Quadratfuß Querschnitt. Man schneidet

nun iteriert Stücke von je 1 Fuß Länge von diesem Balken ab, und legt diese nach einer volumenerhaltenden Verformung um eine Kugel derart herum, daß der Kugel immer neue Schalen hinzugefügt werden, also Schritt für Schritt eine größere Kugel gebildet wird. In dieser Weise wird der ganze Raum ausgefüllt! Denn der Kugeldurchmesser kann bei dieser Prozedur nicht gegen einen endlichen Wert konvergieren, da sich sonst ein endliches Volumen ergeben würde, der Balken aber ein unbegrenztes Volumen hat. Also sind der Balken und der dreidimensionale Raum „deckungsgleich“. Es scheint also vielleicht doch letztendlich „unendlich gleich unendlich“ zu sein, wenn man „gleich“ als „deckungsgleich“ interpretiert?

Albert gibt noch ein zweites Beispiel: Auf die Intervalle $[0, 1/2]$, $[1/2, 3/4]$, $[3/4, 7/8]$, ... legt man abwechselnd weiße und schwarze Steine. Dann nimmt man die schwarzen Steine nacheinander von links nach rechts weg, und rückt dabei die Steine von rechts auf die freiwerdenden Plätze nach. Dadurch wird schließlich jeder Platz mit einem weißen Stein dauerhaft besetzt. Die Menge der weißen Steine ist also deckungsgleich mit der Menge aller verwendeten Steine. All dies steht bei Albert innerhalb eines Kommentars zu „Über den Himmel“ von Aristoteles. Das Balkenbeispiel findet sich genau so auch in einem Kommentar von Oresme, sodaß mutmaßlich einige der von Albert niedergeschriebenen Gedanken der Vierergruppe als Ganzes zuzurechnen sind.

Galileo Galilei (1564 – 1642) ordnet, mutmaßlich unter Kenntnis der Werke von Albert von Sachsen, die Quadratzahlen n^2 den natürlichen Zahlen n bijektiv zu und schließt nachvollziehbar, aber voreilig, daß „groß“, „klein“, „gleich“ offenbar im Unendlichen keinen Sinn machen.

Bernard Bolzano diskutiert in seinen 1848 geschriebenen und 1851 posthum veröffentlichten „Paradoxien des Unendlichen“ ebenfalls die merkwürdige Eigenschaft, daß zwei unendliche Mengen bijektiv aufeinander abbildbar sein können, obwohl die eine eine echte Teilmenge der anderen ist. Aber erst Georg Cantor, der Bolzanos Buch sehr geschätzt hat, erkennt die fundamentale Bedeutung des Mächtigkeitsbegriffs, und geht ihm, motiviert durch einen frühen Erfolg, beharrlich nach. Er ist philosophisch interessiert, und gleichzeitig, aus der mathematischen Praxis der Fourierreihen kommend, an keiner Stelle philosophisch blockiert. Ihn plagen keine Skrupel noch Zweifel, wenn er mit unendlichen Mengen Mathematik nach allen Regeln der Kunst betreibt. Bolzano wollte ebenfalls zeigen, daß die Paradoxien des Unendlichen Scheinparadoxien sind, aber zu vieles bei ihm bleibt Alchemie, er will den Stein der Weisen finden, stolpert dabei über den Kern der Sache, und läßt ihn links liegen. Er erkennt noch nicht, was die Unendlichkeit im Innersten zusammenhält und ordnet. Bolzano bildet so das Endglied der mengentheoretischen Vorzeit.

Die Mengenlehre beginnt irgendwann vor der Erfindung des Zählens und wird von diesem dann wahrscheinlich frühkapitalistisch verdrängt. Sie keimt selbst bei den Griechen nirgendwo auf, im Mittelalter schlägt sie Wurzeln in Kommentaren oft recht dunkler aristotelischer Heiligtümer, und gelangt in dieser Form in die Hände der Neuzeit, wo sie von Wissenschaftlern wie Galilei und Bolzano zwar begrüßt, aber dann doch noch einmal übersehen wird, bis sie schließlich unter Cantor zum ersten Mal ihre Knospen öffnet.

Größenvergleich zweier Mengen: „kleinergleich“

Motiviert durch obige Überlegungen definieren wir:

Definition ($|A| \leq |B|$)

Seien A und B Mengen. Dann ist die *Mächtigkeit von A kleinergleich der Mächtigkeit von B* , in Zeichen $|A| \leq |B|$, falls gilt:

Es existiert ein injektives $f : A \rightarrow B$.

Übung

Sei $|A| \leq |B|$ und $|B| \leq |C|$. Dann gilt $|A| \leq |C|$.

Der folgende Satz hält eine wichtige Äquivalenz zu $|A| \leq |B|$ fest, und gehört zu den mathematischen Resultaten des Typs: „Ich sehe den Beweis.“

Satz

Seien A und B Mengen, und sei $A \neq \emptyset$. Dann sind äquivalent:

- (i) Es existiert $f : A \rightarrow B$ injektiv.
- (ii) Es existiert $g : B \rightarrow A$ surjektiv.

Beweis

(i) $\hookrightarrow$ (ii):

Sei $f : A \rightarrow B$ injektiv. Wir setzen $g' = f^{-1} = \{ (b, a) \mid (a, b) \in f \}$.

Dann ist $g' : \text{rng}(f) \rightarrow A$ surjektiv.

Sei nun $a^* \in A$ beliebig und $g = g' \cup \{ (b, a^*) \mid b \in B - \text{rng}(f) \}$.

Dann ist g eine Funktion von B nach A und $A = \text{rng}(g') \subseteq \text{rng}(g)$, also ist $g : B \rightarrow A$ surjektiv.

(ii) $\hookrightarrow$ (i):

Sei $g : B \rightarrow A$ surjektiv. Für jedes $a \in A$ ist also

$$U_a = g^{-1} \{ a \}$$

nichtleer. Für jedes $a \in A$ fixieren wir ein beliebiges $b_a \in U_a$.

Wir setzen $f = \{ (a, b_a) \mid a \in A \}$. Dann ist $f : A \rightarrow B$ injektiv.

Denn seien $(a_1, b), (a_2, b) \in f$. Dann ist $b \in U_{a_1}$ und $b \in U_{a_2}$,

— also $g(b) = a_1$ und $g(b) = a_2$. Also $a_1 = a_2$.

Der Beweis wiederholt Argumente aus dem Satz über die „Verkettung zur Identität“ aus dem letzten Kapitel. In der Beweisrichtung von (ii) nach (i) findet sich wieder ein Argument vom Typ „ein ...“, wobei wir diesmal die Sprechweise „für jedes ... fixieren wir“ verwenden haben, was mit „für jedes ... wählen wir“ gleichbedeutend ist. (Man spricht in Beweisen gern von „fixieren“ o.ä., wenn ein sonst unbelastetes Zeichen für den Rest des Beweises ein beliebiges Objekt aus einer bestimmten nichtleeren Menge bedeuten soll. Z.B. „Wir fixieren eine Primzahl $p \in \mathbb{N}$, und setzen $A = \{ n \cdot p \mid n \in \mathbb{N} \}$.“)

Eine natürliche Definition eines „größergleich“ für Mengen ist:

Definition ($|B| \geq^* |A|$)

Seien A und B Mengen. Wir setzen

$|B| \geq^* |A|$ falls $A = \emptyset$ oder ein surjektives $f: B \rightarrow A$ existiert.

Ein Ausdruck $|A| \leq^* |B|$ ist dann gleichbedeutend mit $|B| \geq^* |A|$, so wie $|B| \geq |A|$ gleichbedeutend mit $|A| \leq |B|$ ist. Daher der Stern, um die Begriffe erst einmal getrennt zu halten. Aus dem obigen Satz folgt aber, daß für alle Mengen A, B gilt: $|A| \leq |B|$ gdw $|A| \leq^* |B|$.

Der Satz von Cantor-Bernstein (Äquivalenzsatz)

Den Zusammenhang von $|A| = |B|$ und $|A| \leq |B|$ klärt der folgende *Äquivalenzsatz*. Cantor hatte diesen Satz bereits 1883 formuliert. Bewiesen wurde er dann erst in einem von Cantor geleiteten Seminar in Halle im Jahre 1897 von dem damals 19-jährigen Felix Bernstein, der die Arbeiten von Cantor studiert und dabei seinen Beweis entdeckt hatte. Cantor hatte in der Zeit seiner schöpferischen Höchstleistungen in Halle keine Schüler um sich. Die mathematischen Zentren waren damals Göttingen und Berlin. Hätte man dort von Anfang an die Mengenlehre in ihrer Bedeutung erkannt, wären wahrscheinlich Beweise des Äquivalenz- und Vergleichbarkeitssatzes viel früher gefunden worden. Cantor mußte alles alleine machen. Schüler hätten ihn gestützt, nicht nur im Auffinden von Beweisen, sondern auch durch die bloße Möglichkeit zur Diskussion über offene und interessante Fragen der Mengenlehre. Ein einziger kluger Kopf, der Cantor dazu bewegt hätte, sein Wissen um die „Paradoxien der Mengenlehre“ genauer zu formulieren, hätte genügt, um einen Übergang der Mathematik in ihr modernes Zeitalter einzuleiten, ohne daß dabei Porzellan zerschlagen worden wäre. Verlassen wir aber diese, wie der Logiker Georg Kreisel sicherlich sagen würde, „Kaffeehausdiskussion“ und wenden uns dem Beweis des Satzes von Cantor-Bernstein zu.

Der Beweis verwendet die natürlichen Zahlen $\mathbb{N}$ und die rekursive Definition von Objekten: Es wird eine Folge $o_0, o_1, o_2, \dots$ von Objekten konstruiert. Bei der Definition von o_n darf dabei auf alle o_i mit $i < n$ zurückgegriffen werden. [Der Leser betrachte als Beispiel die Definition der Fakultät $n!$ für $n \in \mathbb{N}$: Man setzt $0! = 1$ und definiert dann für $n \in \mathbb{N}$ rekursiv $(n+1)! = n! \cdot (n+1)$.]

So wie man Objekte durch Rekursion definieren kann, kann man Beweise durch Induktion führen: Man will zeigen, daß eine Aussage $A(n)$ für alle $n \in \mathbb{N}$ gilt. Hierzu genügt es zu zeigen: (1) $A(0)$ gilt (*Induktionsanfang*) und (2) Für alle $n \in \mathbb{N}$ gilt: Aus der Gültigkeit von $A(n)$ folgt die Gültigkeit von $A(n+1)$ (*Induktionsschritt* von n nach $n+1$). Eine wichtige Variante des Induktionsschritts ist: (2') Für alle $n \in \mathbb{N}$ gilt: Aus der Gültigkeit von $A(m)$ für alle $m < n$ folgt die Gültigkeit von $A(n)$. Oft reicht die erste Form, zuweilen ist die zweite Form, bei der man im Induktionsschritt alles, was man „bereits bewiesen“ hat, mitführt, von Vorteil. (Zudem liefert (2') den Induktionsanfang (1) kostenlos mit.)

In der Mengenlehre läßt sich beweisen, daß rückbezügliche Definitionen und induktives Beweisen gerechtfertigt sind; für jetzt nehmen wir diese Verfahren als legitim an, wie es auch sonst überall in der Mathematik geschieht. (Vgl. hierzu auch die Diskussion in 2.7.)

Ernst Zermelo hat 1908 den Bernsteinschen Beweis so umgestaltet, daß dabei die natürlichen Zahlen und die Rekursion vermieden werden können. Da die einfache Idee des Beweises, auf die es uns hier ankommt, dabei eher verwischt wird, gehen wir den anderen Weg. Der Äquivalenzsatz zählt wieder zu den Sätzen, deren „Grund der Gültigkeit“ man vor Augen sehen kann, diesmal allerdings wohl nur durch genaues Studium eines Beweises. Der Satz besagt: Aus $|M| \leq |N|$ und $|N| \leq |M|$ folgt $|N| = |M|$, d. h. wenn es Injektionen $f: M \rightarrow N$ und $g: N \rightarrow M$ gibt, so gibt es auch eine Bijektion $h: M \rightarrow N$. Dies ist keineswegs klar.

Wir beweisen vorab einen Satz, dessen Aussage der Intuition sehr entgegen kommt, und aus dem sich der Satz von Cantor-Bernstein in wenigen Zeilen ergibt. Dieser vorbereitende Satz besagt, daß eine zwischen zwei gleichmächtigen Mengen $N \subseteq M$ bzgl. der Inklusion „eingezwickte“ dritte Menge N' ebenfalls zu N und M gleichmächtig ist. Das Gefühl sagt uns, daß dies gelten sollte.

Satz (*Inklusionssatz*)

Seien M, N Mengen mit $N \subseteq M$ und $|N| = |M|$.

Weiter sei N' eine Menge mit $N \subseteq N' \subseteq M$.

Dann gilt $|N'| = |M| [= |N|]$.

Anders formuliert:

Seien A, B, C paarweise disjunkte Mengen mit $|A \cup B \cup C| = |A|$.

Dann gilt auch $|A \cup B \cup C| = |A \cup B|$.

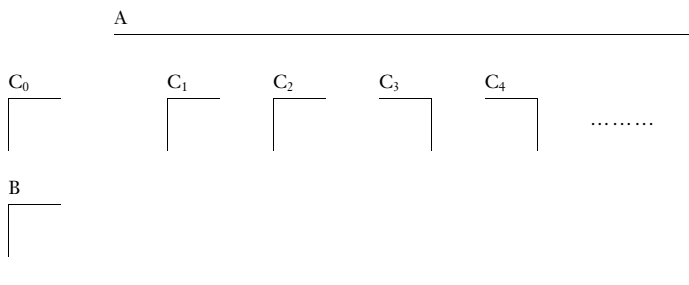
Beweis

Die Äquivalenz der beiden Aussagen ist klar: Wir setzen $A = N$, $B = N' - N$, $C = M - N'$, und in der anderen Richtung $N = A$, $N' = A \cup B$, $M = A \cup B \cup C$. Wir zeigen die A-B-C-Formulierung. Sei also $f: A \cup B \cup C \rightarrow A$ bijektiv. Wir setzen $C_0 = C$ und definieren rekursiv:

$$C_{n+1} = f''C_n \text{ für } n \in \mathbb{N}.$$

Wir bilden also fortwährend die Bilder von C unter der Abbildung f .

Das folgende Diagramm veranschaulicht die Situation:



Es gilt für alle $n \geq 1$: $C_n \subseteq A$ wegen $C_n \subseteq \text{rng}(f) = A$.

Weiter sind die Mengen $C_0, C_1, C_2, \dots$ paarweise disjunkt, wofür die Injektivität von f und $C_0 \cap A = \emptyset$ verantwortlich ist. Obwohl wir diese Aussage für den Beweis nicht brauchen, ist es nützlich, sie sich klar zu machen; zudem rechtfertigt sie obiges Diagramm.

*Beweis der Aussage „für alle $m > n$ ist $C_n \cap C_m = \emptyset$ “
durch Induktion nach n .*

Induktionsanfang $n = 0$:

Für $m \geq 1$ gilt wegen $C_m \subseteq A$, daß $C_0 \cap C_m = \emptyset$.

Induktionsschritt von n nach $n + 1$:

Annahme, es gibt ein

$x \in C_{n+1} \cap C_{m+1} = f''C_n \cap f''C_m$ für ein $m > n$.

Dann ist aber $f^{-1}(x) \in C_n \cap C_m$, im *Widerspruch*
zur Induktionsvoraussetzung.

Also $C_{n+1} \cap C_{m+1} = \emptyset$ für alle $m > n$.

Sei nun

$$C^* = \bigcup_{n \in \mathbb{N}} C_n.$$

Definiere nun eine Funktion $h : A \cup B \cup C \rightarrow A \cup B$ durch:

$$h(x) = \begin{cases} f(x), & \text{falls } x \in C^*, \\ x, & \text{sonst.} \end{cases}$$

Dann ist offenbar h injektiv und $\text{rng}(h) = A \cup B$.

– Also $h : A \cup B \cup C \rightarrow A \cup B$ bijektiv wie gewünscht.

Die Funktion h macht das folgende: Wir bilden C_0 bijektiv durch f auf die Menge $C_1 = f''C_0$ ab. Die Elemente von C_1 schicken wir bijektiv nach C_2 , dann schicken wir die Elemente von C_2 bijektiv nach C_3 usw., wobei wir immer f zum Transport verwenden. Für alle n ist $f : C_n \rightarrow C_{n+1}$ bijektiv und es gilt

$$h|C^* = f|C^* : C^* \rightarrow C^* - C_0$$

bijektiv.

Auf der ganzen Restmenge $(A - \bigcup_{n \geq 1} C_n) \cup B$ ist h die Identität, und damit ist diese Restmenge trivialerweise Teil des Wertebereichs.

Im Grunde ist es ein einfacher Satz.

Bei der Untersuchung der verschiedenen Möglichkeiten, Unendlichkeit zu definieren, werden uns ganz ähnliche Ideen noch einmal begegnen. Wir werden dort die Orbits $x, f(x), f(f(x)), \dots$ von Punkten x unter injektiven Funktionen genauer studieren.

Übung

- (i) Seien $A \subseteq B \subseteq C \subseteq D$, $|A| = |C|$, $|B| = |D|$. Dann gilt $|A| = |D|$.
- (ii) Sei $I = \{x \in \mathbb{R} \mid 0 < x \leq 1\}$, und sei $R = \{x \in I \mid \text{die ersten 100 Nachkommastellen von } x \text{ (in der kanonischen Dezimaldarstellung) sind 0, 1, 2 oder 3}\}$.
Dann gilt $|R| = |I|$.
- (iii) Seien A, B, C Mengen mit $|A| = |A \cup B|$, $|A| = |A \cup C|$
(B und C sind hier nicht notwendig disjunkt).
Dann gilt $|A| = |A \cup B \cup C|$.

Cantor (1883b, § 13): „Hat man irgendeine wohldefinierte Menge M von der zweiten Mächtigkeit, eine Teilmenge M' von M und eine Teilmenge M'' von M' und weiß man, daß die letztere M'' gegenseitig eindeutig abbildbar ist auf die erste M , so ist immer auch die zweite M' gegenseitig eindeutig abbildbar auf die erste und daher auch auf die dritte.

... es scheint mir aber höchst bemerkenswert und ich hebe es daher ausdrücklich hervor, daß dieser Satz *allgemeine* Gültigkeit hat, gleichviel welche Mächtigkeit der Menge M zukommen mag. Darauf will ich in einer späteren Abhandlung näher eingehen und alsdann das eigentümliche Interesse nachweisen, welches sich an diesen allgemeinen Satz knüpft.“

Wir erhalten nun sofort:

Korollar (*Äquivalenzsatz von Cantor-Bernstein*)

Seien M, Q Mengen mit $|M| \leq |Q|$ und $|Q| \leq |M|$.

Dann gilt $|M| = |Q|$.

Beweis

Seien $f : M \rightarrow Q$ injektiv und $g : Q \rightarrow M$ injektiv.

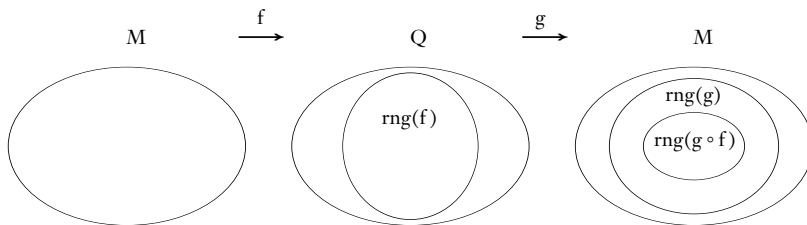
Sei $N = \text{rng}(g \circ f)$, $N' = \text{rng}(g)$.

Dann gilt $N \subseteq N' \subseteq M$ und $|N| = |M|$, denn $g \circ f : M \rightarrow N$ ist bijektiv.

Nach dem Satz oben ist also $|N'| = |M|$.

Aber es gilt $|Q| = |N'|$, denn $g : Q \rightarrow \text{rng}(g) = N'$ ist bijektiv.

– Also $|Q| = |N'| = |M|$.



Mit Hilfe des Satzes kann man die Aufgabe:

(+) Zeige $|A| = |B|$.

in die beiden in vielen Fällen einfacheren Teilaufgaben

(+₁) Zeige $|A| \leq |B|$ und (+₂) Zeige $|B| \leq |A|$.

zerlegen (vgl. $A = B$ gdw $A \subseteq B$ und $B \subseteq A$).

Ernst Zermelos Beweis (veröffentlicht 1908, brieflich mitgeteilt an David Hilbert 1905 und Henri Poincaré (1854 – 1912) 1906) ohne natürliche Zahlen behandeln wir in folgenden Übung. Zermelo beruft sich bei seinem Beweis auf Dedekind. In der Tat hatte Dedekind bereits 1887 einen Beweis des Äquivalenzsatzes gefunden, diesen aber nicht klar herausgestellt – er ist heute im Nachlaßteil seiner „Gesammelten Werke“ zu finden [Dedekind 1930 – 1932, Band III, S. 447ff]. Unabhängig wurde dieser Beweis auch entdeckt von Peano 1906 und Alwin Korselt (1864 – 1947) 1911.

Übung

Wir zeigen den Inklusionssatz, wie oben folgt dann der Äquivalenzsatz.

Sei $f : A \cup B \cup C \rightarrow A$ bijektiv, wobei A, B, C paarweise disjunkt.

Wir setzen $Z = \bigcap \{ D \subseteq A \cup C \mid C \subseteq D \text{ und für alle } x \in D \text{ ist } f(x) \in D \}$.

- (i) Es gilt $C \subseteq Z$ und für alle $x \in Z$ ist $f(x) \in Z$.
- (ii) Ist $x \in Z - C$, so ist $f^{-1}(x) \in Z$. Folglich $f''Z = Z - C$.
- (iii) Definiere $h : A \cup B \cup C \rightarrow A \cup B$ durch
 $h(x) = f(x)$, falls $x \in Z$, $h(x) = x$, falls $x \notin Z$.
 Dann ist $h : A \cup B \cup C \rightarrow A \cup B$ bijektiv.

[De facto gilt $Z = C^*$ mit C^* wie im Beweis oben und die konstruierten Bijektionen h sind dieselben. Z wird hier aber „von oben“ als Schnitt definiert; der Aufbau von $Z = C^*$ „von unten“ (induktiv) gibt sicher ein klareres Bild von Z .]

Wir geben nun noch einen weiteren Beweis des Satzes von Cantor-Bernstein, der 1906 von Julius König (1849 – 1913) gefunden wurde. Im Grunde konstruieren wir die gleiche Bijektion wie oben, aber der Beweis erzeugt eine etwas andere Vorstellung von dieser Bijektion. Der Inklusionssatz wird nicht verwendet.

Weiterer Beweis des Satzes von Cantor-Bernstein

Seien $f : M \rightarrow Q$ und $g : Q \rightarrow M$ injektiv. O. E. seien M und Q disjunkt.

[Andernfalls ersetze M durch $M \times \{0\}$ und Q durch $Q \times \{1\}$ und modifiziere f und g entsprechend; ein bijektives $h' : M \times \{0\} \rightarrow Q \times \{1\}$ liefert dann auch ein bijektives $h : M \rightarrow Q$. Der Beweis funktioniert auch direkt für beliebige M und Q . Aber für das innere Auge (und für Skizzen) ist der Fall $M \cap Q = \emptyset$ ästhetischer.]

Wir suchen wieder eine Bijektion $h : M \rightarrow Q$.

Wir klassifizieren hierzu die Elemente x aus M nach dem Typ der durch sie laufenden f - g -Urbildkette: Sei $x \in M$. Wir definieren solange wie möglich:

$$x = x_0, \quad x_1 = g^{-1}(x_0), \quad x_2 = f^{-1}(g^{-1}(x_0)) = f^{-1}(x_1), \quad x_3 = g^{-1}(x_2), \quad x_4 = f^{-1}(x_3), \dots$$

Allgemein setzen wir

$$x_{n+1} = \begin{cases} g^{-1}(x_n), & \text{falls } n \text{ gerade,} \\ f^{-1}(x_n), & \text{falls } n \text{ ungerade,} \end{cases}$$

solange die Urbilder existieren!

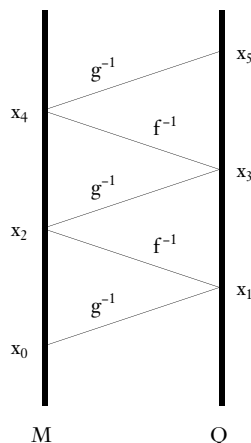
Ist x_n für alle $n \in \mathbb{N}$ definiert, so sagen wir: x ist vom Typ I.

Ist das letzte definierte $x_n \in M$, d. h. $g^{-1}(x_n)$ existiert nicht, so sagen wir: x ist vom Typ II.

Ist das letzte definierte $x_n \in Q$, d. h. $f^{-1}(x_n)$ existiert nicht, so sagen wir: x ist vom Typ III.

Offenbar kann ein x nicht von zwei verschiedenen Typen zugleich sein.

Der Leser verfolge an einer Skizze den Ziehharmonika-Kurs $x = x_0, x_1, x_2, \dots$ der Urbilder von x unter den Abbildungen f und g . Dieser Kurs kann abbrechen – u. U. schon bei x_0 –, weil f und g nicht notwendig surjektiv sind. Er kann unendlich sein, z. B. für den Fall $g(f(x)) = x$; hier ist die Urbildfolge $x, f(x), x, f(x), \dots$. Eine Urbildfolge eines x vom Typ I kann natürlich auch aus unendlich vielen paarweise verschiedenen Elementen bestehen.



Wir setzen nun für $x \in M$:

$$h(x) = \begin{cases} f(x), & \text{falls } x \text{ vom Typ I oder II ist,} \\ g^{-1}(x), & \text{falls } x \text{ vom Typ III ist.} \end{cases}$$

- Dann ist $h : M \rightarrow Q$ die gewünschte Bijektion, wovon der Leser sich mit Freuden überzeugt.

Auch $h : M \rightarrow Q$ mit $h(x) = f(x)$, falls x vom Typ II, $h(x) = g^{-1}(x)$, falls x vom Typ I oder III wäre eine Bijektion von M nach Q . Auf Typ I könnte man weiter auch doppelt verschieben: $h(x) = f(f(x))$, usw. Wichtig ist nur, daß wir die schließlich endenden Urbildketten gemäß ihrem Ende in M oder Q unterschiedlich behandeln.

Schließlich kann man den Satz von Cantor-Bernstein auch mit einem auf Bronisław Knaster und Alfred Tarski zurückgehenden Fixpunktsatz beweisen, wobei hier der Fixpunktsatz interessanter ist als seine Anwendung auf den Äquivalenzsatz, die im Grunde nur den Beweis von Dedekind-Zermelo wiederholt. Wir behandeln diesen Ansatz in der folgenden Übung.

Übung

Sei M eine Menge. Eine Funktion $g : \mathcal{P}(M) \rightarrow \mathcal{P}(M)$ heißt *monoton*, falls für alle $x, y \in \mathcal{P}(M)$ gilt: $x \subseteq y$ folgt $g(x) \subseteq g(y)$.
 x ist ein Fixpunkt von g , falls $g(x) = x$ gilt.

- (i) Seien M eine Menge und $g : \mathcal{P}(M) \rightarrow \mathcal{P}(M)$ monoton.

Dann existiert ein kleinster Fixpunkt von f , d. h. es gibt ein $y \in \mathcal{P}(M)$ mit: $g(y) = y$ und für alle $x \subseteq y$ gilt $g(x) \neq x$.

[Wir setzen $y = \bigcap \{x \subseteq M \mid g(x) \subseteq x\}$ (es gilt z. B. $g(M) \subseteq M$).]

- (ii) Zeigen Sie den Inklusionssatz (und damit den Satz von Cantor-Bernstein) mit Hilfe eines kleinsten Fixpunktes einer geeignet definierten monotonen Funktion.

[Sei $D = A \cup B \cup C$ und $f : D \rightarrow A$ bijektiv.

Wir definieren $g : \mathcal{P}(D) \rightarrow \mathcal{P}(D)$ durch $g(x) = f''x \cup C$. Dann ist g monoton. Der kleinste Fixpunkt von g ist gerade unser früheres C^* .]

Analog existiert für ein monotones $g : \mathcal{P}(M) \rightarrow \mathcal{P}(M)$ ein $\subseteq$ -größter Fixpunkt. Für die monotone Funktion g in (ii) ist $(A \cup B \cup C) - B^*$ der größte Fixpunkt von g , wobei B^* analog zu C^* definiert ist, also $B = \bigcup_{n \in \mathbb{N}} B_n$ mit $B_0 = B$, $B_{n+1} = f''B_n$. Weiter ist die Funktion $h' : D \rightarrow A \cup B$ mit $h'(x) = f(x)$ für $x \notin B^*$, $h'(x) = x$ für $x \in B^*$ bijektiv. Hier ist also der Anteil der Identität kleiner, falls kleinster und größter Fixpunkt verschieden sind, d. h. falls $B^* \cup C^* \neq D$. Alle bekannten „elementaren“ Beweise des Satzes von Cantor-Bernstein konstruieren letztendlich eine der Abbildungen h oder h' .

Die obigen Beweise des Satzes von Cantor-Bernstein sind zwar trickreich, insgesamt aber weniger abstrakt als etwa der sehr übersichtliche und einfache Beweis, daß aus der Existenz einer Surjektion $f : A \rightarrow B$ die Existenz einer Injektion $g : B \rightarrow A$ folgt. Ein „ein ...“ wird hier nirgendwo gebraucht. Während sonst, wie wir sehen werden, im Bereich der Mächtigkeitstheorie Vermutungen mit den schwersten Geschützen solange angeschossen werden, bis sie sich endlich beweisen oder widerlegen lassen, genügt zum Beweis des Satzes von Cantor-Bernstein eine letztendlich simple trojanische Kriegslist.

Wir verweisen den Leser auf [Deiser 2008] und [Rautenberg 1987] für weitere historische Bemerkungen und Beweise des Satzes.

Strikt kleinere Mächtigkeiten

Definition ($|A| < |B|$)

Seien A, B Mengen. A ist von (echt) kleinerer Mächtigkeit als B , in Zeichen $|A| < |B|$, falls $|A| \leq |B|$ und $|A| \neq |B|$.

Also $|A| < |B|$ gdw „es existiert eine Injektion von A nach B , aber keine Bijektion von A nach B “.

Für die Transitivität von $|A| < |B|$ wird der Satz von Cantor-Bernstein benötigt:

Übung

$|A| < |B|$ und $|B| < |C|$ folgt $|A| < |C|$.

[Der Satz von Cantor-Bernstein wird verwendet, um den Fall $|A| = |C|$ auszuschließen.]

In der älteren Literatur wurde das „kleiner“ bei Mächtigkeiten oft etwas anders eingeführt. Für Mengen A, B definiert man hier $|A| <_2 |B|$, falls eine Injektion von A nach B , aber keine Injektion von B nach A existiert, d. h. man setzt $|A| <_2 |B|$, falls $|A| \leq |B|$ aber $\text{non}(|B| \leq |A|)$. $<_2$ ist offenbar irreflexiv. Weiter ist $<_2$ transitiv, und hierfür wird der Satz von Cantor-Bernstein nicht benutzt (Beweis von „ $|A| < |B| \leq |C|$ folgt $|A| <_2 |C|$ “ ohne Cantor-Bernstein als Übung). Der Satz von Cantor-Bernstein wird benutzt um zu zeigen, daß $<$ und $<_2$ identisch sind.

Was ist $|M|$?

Für eine Menge M hat $|M|$ noch keine Bedeutung an sich, streng definiert sind lediglich die Beziehungen $|M| = |N|$, $|M| \leq |N|$ und $|M| < |N|$. Durch den ständigen Gebrauch von Wendungen wie „die Mächtigkeit von M ist kleiner als die Mächtigkeit von N “ wird man aber irgendwann „Mächtigkeit“ als ein selbständiges Wort ansehen, bei dem nicht nur in bezug auf zwei begleitende Mengen auch ein Begriff ist. Es ist sehr schwer, diesen Begriff der Mächtigkeit einer Menge genau zu definieren, und es ist nicht ohne Witz, daß gerade die Mengenlehre als Sprache, die sonst mit Leichtigkeit Konzepte aus der Mathematik entgegennimmt und fehlerfrei gemeißelte Definitionen zurückgibt, mit einem der ureigensten Begriffe der Mengenlehre als Theorie größere handwerkliche Schwierigkeiten hat. Bis weit in die 20er Jahre des 20. Jahrhunderts hat es gedauert, bis man eine wirklich befriedigende Definition von $|M|$ zusammenschweißt hatte.

Unser Vorgehen ist nun dieses: Bis zu Kapitel 12 werden wir $|M|$ nicht als Objekt betrachten, sondern rein relational mit $|M| = |N|$, $|M| \leq |N|$ und $|M| < |N|$ auskommen. Es trifft sich glücklich, daß nicht nur die Resultate, sondern auch die Wege zu ihrer Entdeckung sich in dieser Form sehr treu wiedergeben lassen: im Vordergrund steht die Konstruktion von Injektionen und Bijektionen *zwischen* zwei Mengen M und N , nicht die Manipulation *eines* Objektes $\alpha = |M|$. Kapitel 12 führt dann in die Arithmetik mit Mächtigkeiten (oder gleichbedeutend: Kardinalzahlen) ein, und bis dahin wird hoffentlich jeder Leser unter „Mächtigkeit von M “ etwas verstehen, und einen nicht zuletzt notationell flexiblen algebraischen Kalkül mit Mächtigkeiten als befreiend empfinden. Es wird klar sein, daß sich alle neuen Resultate in die alte relationale Sprache zurückübersetzen ließen, wenn man denn dies wirklich wollte. Wir nehmen dann also mit „ $|M|$ als Objekt“ zwar in der Tat eine kleine definitorische Schuld auf – denn streng definieren können wir $|M|$ in Kapitel 12 immer noch nicht –, bleiben aber in der Lage, sie jederzeit zurückzuzahlen. Die Hypothek erweist sich zudem als Investition: Der Kalkül reproduziert mühelos alte Beweise, er vereinfacht die Theorie und suggeriert neue Fragestellungen. Was will man mehr.

Man will mehr. Man will eine saubere Definition. Diese entwickeln wir im zweiten Abschnitt. Zunächst geben wir dort eine taktische Definition von Kardinalzahl über eine neue Hypothek, nämlich die der Ordinalzahl, bis wir zuletzt „Ordinalzahl“ formal befriedigend definieren können, und damit jede Hypothek zurückzahlen. $|M|$, die Mächtigkeit oder Kardinalität von M , wird schließlich exakt definiert sein als eine *Zahl*, wobei hier ein erweiterter Zahlbegriff zugrundeliegt. Die Reihe der natürlichen Zahlen wird durch die transfiniten Zahlen fortgesetzt werden, und im weiten Raum der Unendlichkeit finden sich dann genügend ausgezeichnete und sichere Bojen, mit denen der Begriff der Mächtigkeit sicher verankert werden kann.

Würden diese definitorischen Schwierigkeiten dem Verständnis des Stoffes irgendwie entgegenstehen, so hätte der Autor keine Sekunde gezögert, das Pferd

von hinten aufzuzäumen. Aber im Gegenteil scheint die Freundschaft zu abstrakten Begriffen durch persönliche Begegnungen weit besser gefördert zu werden als durch eine notarielle Beglaubigung ihres gesetzestreuen Verhaltens. Später wird man diese dann zu schätzen wissen, zumal sie sich in bestechend schöner Handschrift anbietet.

Georg Cantor hat die Mächtigkeit $|M|$ einer Menge M als das Gemeinsame aller Mengen definiert, die gleichmächtig mit M sind, d.h. für die eine Bijektion zu M existiert. Dies genügte ihm zeitlebens als eine befriedigende Definition.

Georg Cantor (1887): „Unter Mächtigkeit oder Kardinalzahl einer Menge M (die aus wohlunterschiedenen, begrifflich getrennten Elementen $m, m', \dots$ besteht und insofern bestimmt und abgegrenzt ist) verstehe ich den Allgemeinbegriff oder Gattungsbegriff (universale), welchen man erhält, indem man bei der Menge sowohl von der Beschaffenheit ihrer Elemente, wie auch von allen Beziehungen, welche die Elemente, sei es unter einander, sei es zu anderen Dingen haben, also im besonderen auch von der Ordnung, welche unter den Elementen herrschen mag, abstrahiert und nur auf das reflektiert, was allen Mengen gemeinsam ist, die mit M äquivalent sind. Ich nenne aber zwei Mengen M und N äquivalent, wenn sie sich gegenseitig eindeutig Element für Element einander zuordnen lassen ...“

Dabei ist dieses Gemeinsame durchaus als ein Objekt zu denken, nicht als eine Art kontemplativer Zusammenschau aller charakteristischen Merkmale gleichmächtiger Mengen oder – fast noch schlimmer – extensional als Zusammenfassung aller Mengen N mit $|N| = |M|$; derartige N gibt es uferlos viele: für jedes x hat $\{x\}$ genau ein Element und damit wäre bereits die Mächtigkeit 1 als Zusammenfassung aller einelementigen Mengen fast schon satirisch definiert. (Zudem geraten wir durch solche unüberschaubaren Zusammenfassungen in die Schwierigkeiten des naiven Komprehensionsprinzips.) Die doppelte Abstraktion $|M|$ bei Cantor hinterläßt ein Objekt, das möglichst frei von speziellen Eigenschaften ist. Wir diskutieren im zweiten Abschnitt Cantors Vision einer Kardinalzahl noch einmal genauer.

Abraham Fraenkel hat in seiner „Einleitung in die Mengenlehre“ auf die Probleme, aber auch auf den Wert der intuitiven Begriffsbildung durch Abstraktion hingewiesen, und mit seinen Worten beenden wir vorläufig diese Diskussion – und dieses Kapitel.

Abraham Fraenkel über Begriffsbildungen durch Abstraktion

„Gegen die vorstehend angegebene Art der Einführung der unendlichen Kardinalzahlen läßt sich allerdings der Einwand erheben, daß die Kennzeichnung ‚das Gemeinsame aller (oder je zweier) untereinander äquivalenter unendlicher Mengen‘ keine scharfe Begriffsbildung ist, wie man sie von einer mathematischen Definition mit Recht verlangen muß...“

Eine tiefere Einsicht in das hier vorliegende Problem gewinnt man durch die Erkenntnis, daß eine durchaus entsprechende Sachlage zu vielen wichtigen Begriffsbildungen

innerhalb (und gelegentlich auch außerhalb) der Mathematik geführt hat. So gelangt man, indem man immer die gemeinsame ‚typische Eigenschaft‘ aller Individuen einer ganzen Klasse aufsucht, z. B. von jeder Klasse aller untereinander *parallelen* Strahlen zum Begriff der ihnen gemeinsamen *Richtung* ..., von jeder Gesamtheit aller untereinander *ähnlichen* Figuren zum Begriff der ihnen gemeinsamen *Gestalt*, usw. Die typische Eigenschaft ordnet all die Objekte, denen sie zukommt, jeweils in eine Klasse ein, in die Klasse der Objekte von ‚gleichem Typus‘ ... Dann läßt sich nämlich immer einer beliebigen Klasse von Objekten desselben Typus ein Begriff zuordnen, der der *Typus* eines jeden Objektes der Klasse ... heißen soll; der Typus von Mengen hinsichtlich der Äquivalenz [Gleichmächtigkeit] ist so die Kardinalzahl, ebenso wie die Richtung den Typus orientierter gerader Linien hinsichtlich des Parallelismus darstellt. Man drückt sich vielfach so aus: Der Richtungsbegriff entsteht aus dem Begriff der orientierten Geraden, indem man von der Lage der Geraden im Raum abstrahiert; ebenso der Kardinalzahlbegriff aus dem Begriff der Menge, indem man von der Natur (sowie einer etwaigen Anordnung) der Elemente absieht...”

(Abraham Fraenkel 1928, „*Einleitung in die Mengenlehre*“)

5. Der Vergleichbarkeitssatz

Der Vergleichbarkeitssatz, von Cantor lange Zeit vermutet, wurde zum ersten Mal streng bewiesen durch Ernst Zermelo (1904; der Satz ist ein Korollar zum Zermeloschen Wohlordnungssatz). Er beantwortet die an dieser Stelle der Diskussion wohl natürlichste Frage positiv, nämlich: Sind je zwei Mengen in ihrer Mächtigkeit vergleichbar?

Der Beweis des Satzes ist der härteste Brocken dieser Einführung, und der nicht allzu ehrgeizige Leser kann zunächst nur seine Aussage zur Kenntnis nehmen und den Beweis überschlagen.

Satz (Vergleichbarkeitssatz)

Seien M, N Mengen. Dann gilt $|M| \leq |N|$ oder $|N| \leq |M|$.

Beweis (Technik von Ernst Zermelo und Erhard Schmidt, 1904 und 1908)

Sei $\mathcal{X} = \{f \mid \text{dom}(f) = M' \subseteq M, \text{rng}(f) = N' \subseteq N, f: M' \rightarrow N' \text{ bijektiv}\}$.

Es genügt zu zeigen:

($\diamond$) Es existiert ein $f \in \mathcal{X}$ mit $\text{dom}(f) = M$ oder $\text{rng}(f) = N$.

Im Fall $\text{dom}(f) = M$ ist $|M| \leq |N|$, denn dann ist $f: M \rightarrow N$ injektiv, und im Fall $\text{rng}(f) = N$ ist $|N| \leq |M|$, denn dann ist $f^{-1}: N \rightarrow M$ injektiv.

Die Menge $\mathcal{X}$ besteht aus Approximationen an das erwünschte Objekt, und wir müssen zeigen, daß ein erwünschtes Objekt f wie in ($\diamond$) tatsächlich in $\mathcal{X}$ vorkommt.

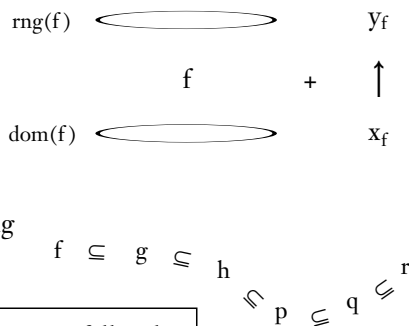
Annahme, ein f wie in ($\diamond$) existiert nicht.

Für jedes $f \in \mathcal{X}$ existieren dann
 $x_f \in M$ und $y_f \in N$ mit $x_f \notin \text{dom}(f)$
 und $y_f \notin \text{rng}(f)$.

Für $f \in \mathcal{X}$ sei

$f^* = f \cup \{(x_f, y_f)\}$.

Dann ist $f^* \in \mathcal{X}$ eine echte Fortsetzung
 von f um genau ein Element.



Ein $T \subseteq \mathcal{X}$ heißt eine (*nichttriviale*) Kette, falls gilt:
 $T \neq \emptyset$ und für alle $f, g \in T$ gilt $f \subseteq g$ oder $g \subseteq f$.

Beweis von ($\clubsuit$)

Sei $T = \{ g \in T^* \mid g \subseteq f \text{ oder } f^* \subseteq g \}$.

Dann ist T geschlossen:

zu (i): Es gilt $f_0 \subseteq f$, also $f_0 \in T$.

zu (ii): Sei $g \in T$. Dann gilt $g \subseteq f$ oder $f^* \subseteq g$.

Wegen f Schnitt gilt zudem $g^* \subseteq f$ oder $f \subseteq g$.

Ist $f^* \subseteq g \subseteq g^*$ oder $g^* \subseteq f$, so ist offenbar $g^* \in T$.

Andernfalls ist $g \subseteq f$ und $f \subseteq g$, also $f = g$.

Also $f^* \subseteq g^*$, und damit $g^* \in T$.

zu (iii): Sei $T' \subseteq T$ eine Kette, und sei $h = \bigcup T'$.

Wegen $T' \subseteq T$ gilt $g \subseteq f$ oder $f^* \subseteq g$ für alle $g \in T'$.

Ist $g \subseteq f$ für alle $g \in T'$, so ist $h = \bigcup T' \subseteq f$, also $h \in T$.

Ist $f^* \subseteq g$ für ein $g \in T'$, so ist $f^* \subseteq h$, also $h \in T$.

Also gilt $T = T^*$, und dies genügt für ($\clubsuit$), da f ein Schnitt ist.

Wir setzen nun:

$T = \{ f \in T^* \mid f \text{ ist ein Schnitt von } T^* \}$.

T ist geschlossen:

zu (i):

f_0 ist ein Schnitt von T^* nach (+).

zu (ii):

Sei f ein Schnitt von T^* .

Dann gilt nach ($\clubsuit$) für alle $g \neq f, g \in T^*$:

$g^* \subseteq f$ oder $f^* \subseteq g$.

Dann gilt aber für alle $g \in T^*$ wegen $f \subseteq f^*$:

$g^* \subseteq f^*$ oder $f^* \subseteq g$.

(Dies ist trivial für $g = f$.)

Also ist f^* ein Schnitt von T^* , und damit $f^* \in T$.

zu (iii):

Sei $T' \subseteq T$ eine Kette, und sei $h = \bigcup T'$.

Sei $g \in T^*$. Wir zeigen: $g^* \subseteq h$ oder $h \subseteq g$.

Dann ist h ein Schnitt, und somit $h \in T$.

Da jedes $f \in T'$ ein Schnitt ist, gilt $g^* \subseteq f$ oder $f \subseteq g$ für alle $f \in T'$.

Ist $f \subseteq g$ für alle $f \in T'$, so gilt $h = \bigcup T' \subseteq g$.

Ist $g^* \subseteq f$ für ein $f \in T'$, so gilt $g^* \subseteq h$ wegen $f \subseteq h$ für jedes $f \in T'$.

Also ist T geschlossen, und damit ist $T = T^*$.

Also gilt $g^* \subseteq f$ oder $f \subseteq g$ für alle $f, g \in T^*$,

insbesondere ist also T^* eine Kette.

Sei $f = \bigcup T^*$. Dann ist nach (iii) $f \in T^*$, also nach (ii) $f^* \in T^*$.

Also $\bigcup T^* = f \subset f^*$ und $f^* \in T^*$. *Widerspruch!*

– Also ist die Annahme falsch und es gilt ($\diamond$), d.h. es existiert ein $f \in \mathcal{L}$ wie gewünscht.

Im Beweis wird T^* als der Schnitt über die geschlossenen Teilmengen $T \subseteq \mathcal{X}$ definiert. Im folgenden wird dann mehrfach benutzt, daß T^* keinen überflüssigen Ballast mehr enthält. T^* besteht nur noch aus den Elementen von $\mathcal{X}$, die eine geschlossene Menge einfach haben muß, um geschlossen zu sein. Alles andere fällt bei der Schnittbildung weg, so etwa der Bereich C im Diagramm oben. In geschlossenen Mengen muß zwischen f und f^* nichts liegen, und T^* verzichtet als spartanische geschlossene Menge auf jeden Schnickschnack.

Der Beweis enthält wieder eine Definition von Objekten der Form „ein ...“, und zwar bei der Definition von x_f und y_f , und damit bei der Definition von f^* ; zu Beginn des Beweises wählen wir für jedes $f \in \mathcal{X}$ Zeugen x_f und y_f aus jeweils nichtleeren von f abhängigen Mengen. Den „ $*$ “ können wir als eine Funktion auf $\mathcal{X}$ auffassen, die ein $f \in \mathcal{X}$ auf $f^* \in \mathcal{X}$ abbildet.

Die logisch einwandfreie uniforme Definition von f^* für alle f gleich zu Beginn – nicht etwa eine schrittweise Erweiterung von Objekten f „während“ des Beweises –, ist die Idee, die Erhard Schmidt (1876 – 1959) zum Beweis des Satzes beigetragen hat, wie Zermelo in seiner Arbeit von 1904, wo die Technik zum ersten Mal angewendet wird, ausdrücklich festhält (vgl. 2.5). Der Rest ist, einschließlich der Betonung der Auswahlakte an sich, die „Zutat“ von Ernst Zermelo.

Der Vergleichbarkeitssatz zeigt den linearen Charakter der Mächtigkeiten und ist Grundvoraussetzung dafür, daß wir – in einem später zu präzisierenden Sinn – $|M|$ als Anzahl, genauer als sogenannte Kardinalzahl ansehen können. Was auch immer wir unter Anzahl verstehen wollen, so ist es doch sicher dieses, daß von zwei Anzahlen eine größergleich der anderen ist.

Der obige Beweis ist vom Typ: Es existiert ein maximales Objekt für eine bestimmte Eigenschaft. In unserem Fall: ein nicht mehr verlängerbares $f \in \mathcal{X}$. Später werden wir allgemeine und in der Mathematik vielseitig einsetzbare Sätze über die Existenz maximaler Objekte zeigen, das Hausdorffsche Maximalitätsprinzip und das sogenannte Zornsche Lemma, aus denen sich der Vergleichbarkeitssatz dann leicht ableiten läßt.

Wir halten noch ein Korollar fest, das wir aus der freien Wahl von $f_0 \in \mathcal{X}$ gewinnen. Wir können $f_0 \in \mathcal{X}$ bereits zu Beginn des Beweises fixieren und dann überall mit der Menge $\mathcal{X}_0 = \{f \in \mathcal{X} \mid f_0 \subseteq f\}$ anstelle von $\mathcal{X}$ arbeiten. Dann zeigt der Beweis:

Korollar

Seien M, N Mengen, $M' \subseteq M$, $N' \subseteq N$, $f_0 : M' \rightarrow N'$ bijektiv.

Dann existiert eine injektive Fortsetzung f von f_0 mit:

$\text{dom}(f) = M$, $\text{rng}(f) \subseteq N$ oder $\text{dom}(f) \subseteq M$, $\text{rng}(f) = N$.

Ist $|M| < |N|$, so existiert eine injektive Fortsetzung f von f_0 mit $f : M \rightarrow N$.

Es gibt also keine Approximationen f_0 , die in eine Sackgasse führen würden: Jede solche Approximation läßt sich bis ans Ziel bringen.

Übung

Gilt in der Situation des Korollars lediglich $|M| \leq |N|$, so existiert im allgemeinen keine injektive Fortsetzung f von f_0 mit $f : M \rightarrow N$.

Zermelosysteme und ein allgemeines Prinzip

Beim Betrachten des Beweises des Vergleichbarkeitssatzes fällt auf, daß ab der Definition einer geschlossenen Menge von der speziellen Natur der Elemente f und ihren Erweiterungen f^* überhaupt nicht mehr die Rede ist. Wir verwenden nicht einmal, daß f^* eine einelementige Erweiterung von f ist (was den Beweis minimal vereinfachen würde). Es ist nun leicht, ein allgemeines Prinzip aus dem Beweis herauszufiltern, das des Pudels Kern zum Vorschein bringt, oder, anders formuliert, das den Hauptteil des Beweises in den Rang eines mathematischen Satzes erhebt. Dies ermöglicht uns, in ähnlichen Fällen alles griffbereit zu haben, ohne das ganze Argument immer wieder neu aus dem Keller holen zu müssen.

Offiziell definieren wir den Begriff einer Kette:

Definition ($\subseteq$ -Kette)

Eine Menge T heißt eine $\subseteq$ -Kette oder kurz *Kette*, falls gilt:

Für alle $x, y \in T$ gilt $x \subseteq y$ oder $y \subseteq x$.

Eine Kette T heißt *Kette in einer Menge A* , falls $x \subseteq A$ für alle $x \in T$.

Die leere Menge gilt nun, aus Gründen der späteren Fügsamkeit des Begriffs, als Kette.

Eine wichtige Eigenschaft, die im Beweis des Vergleichbarkeitssatzes auftrat, war die Abgeschlossenheit von $\mathcal{Z}$ unter der Vereinigung jeder nichttrivialen Kette:

Definition (Zermelosysteme und ihre Ziele)

Sei $\mathcal{Z}$ eine nichtleere Menge. $\mathcal{Z}$ heißt ein *Zermelosystem*, falls gilt:

Ist $T \subseteq \mathcal{Z}$ eine nichtleere Kette, so ist $\bigcup T \in \mathcal{Z}$.

$x \in \mathcal{Z}$ heißt ein *Ziel* des Zermelosystems $\mathcal{Z}$, falls gilt:

Es gibt kein $y \in \mathcal{Z}$ mit $x \subset y$.

In vielen Fällen – so etwa für $\mathcal{Z}$ aus dem obigen Beweis – wird $\emptyset \in \mathcal{Z}$ gelten, und dann ist $\bigcup T \in \mathcal{Z}$ sogar für alle Ketten T , denn $\bigcup \emptyset = \emptyset$.

Satz (Satz über Zermelosysteme)

Sei $\mathcal{Z}$ ein Zermelosystem, und sei $x_0 \in \mathcal{Z}$.

Dann existiert ein Ziel $x \in \mathcal{Z}$ mit $x_0 \subseteq x$.

Der Leser kann sich gemütlich zurücklehnen. Der Beweis ist genau der obige – nur daß jedes f jetzt ein x ist, da nicht suggeriert werden soll, daß $\mathcal{Z}$ aus Funktionen bestehen muß.

Beweis

Annahme, ein solches Ziel existiert nicht. Wir definieren dann eine Schmidt-Expansion $*$ für das Zermelo-System $\mathcal{Z}_0 = \{ x \in \mathcal{Z} \mid x_0 \subseteq x \}$ durch Auswahlakte. Für jedes $x \in \mathcal{Z}_0$ sei hierzu

$x^* = \text{„ein } y \in \mathcal{Z} \text{ mit } x \subset y\text{“}.$

Ein $T \subseteq \mathcal{L}_0$ heißt *geschlossen*, falls gilt:

- (i) $x_0 \in T$.
- (ii) Für alle $x \in \mathcal{L}_0$ gilt: $x \in T$ *folgt* $x^* \in T$.
- (iii) Ist $T' \subseteq T$ eine nichtleere Kette, so ist $\bigcup T' \in T$.

$\mathcal{L}_0$ selbst ist als Zermelosystem offenbar geschlossen.

Wir definieren die Zermelo-Reduzierung geschlossener Mengen durch:

$$T^* = \bigcap \{ T \subseteq \mathcal{L}_0 \mid T \text{ ist geschlossen} \}.$$

Dann ist T^* geschlossen, und es gilt $x_0 \in T^*$.

Genau wie im Beweis des Vergleichbarkeitssatzes folgt:

T^* ist eine Kette.

Sei $x = \bigcup T^*$. Dann ist nach (iii) $x \in T^*$, also nach (ii) $x^* \in T^*$.

– Also $\bigcup T^* = x \subset x^*$ und $x^* \in T^*$. *Widerspruch!*

Wir werden diesen Satz in Kapitel 12 mehrfach verwenden, und im zweiten Abschnitt eine ordnungstheoretische Umformulierung des Satzes über die Existenz von Zielen in Zermelosystemen kennenlernen, das „Zornsche Lemma“ – formuliert von Max Zorn (1906 – 1993) als Schweizer Taschenmesser 1935 –, und daneben das nahverwandte Hausdorffsche Maximalitätsprinzip (1914).

Nimmt man den Satz über Zermelosysteme als „Blackbox“, so ergibt sich leicht ein Beweis des Vergleichbarkeitssatzes:

Übung

Beweisen Sie den Vergleichbarkeitssatz mit Hilfe des Satzes über Zermelosysteme.

[Seien $M, N, \mathcal{L}$ wie im Vergleichbarkeitssatz. Ist $T \subseteq \mathcal{L}$ eine Kette, so ist $\bigcup T \in \mathcal{L}$. Ein Ziel von $\mathcal{L}$ liefert $|M| \leq |N|$ oder $|N| \leq |M|$.]

Unsere Intuition über die Gültigkeit des Satzes

Wir besprechen noch ein intuitives Argument für die Richtigkeit des Vergleichbarkeitssatzes, das sich an unserem Algorithmus des Abtragens orientiert, der ja für je zwei Nußhaufen die Vergleichbarkeit zeigt. Seien hierzu M und N zwei beliebige Mengen. Wie früher bilden wir wiederholt Paare (x, y) mit $x \in M$, $y \in N$, und entfernen dabei x und y aus M und N bzw. ihren verbliebenen Resten. Auch nach unendlich vielen Paarbildungen können die Restmengen von M und N immer noch Elemente enthalten. Dies hindert uns aber nicht daran, das Verfahren mit diesen nach unendlich vielen Schritten verbliebenen Restmengen fortzusetzen. Und dies tun wir solange, bis eine der beiden Mengen aufgebraucht ist. Die gebildeten Paare (x, y) bilden nun eine Injektion von M nach N , falls M aufgebraucht ist, oder ihre Umkehrungen $(x, y)^{-1} = (y, x)$ bilden eine Injektion von N nach M , falls N aufgebraucht ist.

Man kann nun obigen Beweis des Vergleichbarkeitssatzes als eine strenge mathematische Durchführung dieser Idee ansehen. Die Bildung von f^* etwa beschreibt die Entfernung eines weiteren Paares, falls f alle bislang entfernten Paare bezeichnet. Die dritte Bedingung einer geschlossenen Menge beschreibt unser Fortfahren nach unendlich vielen Schritten, „im Limes“. T^* ist dann schließlich die Menge all der Mengen von abgetragenen Paaren, die wir erhalten, wenn wir mit f_0 – etwa $f_0 = \emptyset$ – beginnen, von f jeweils zu f^* übergehen, und im Limes jeweils die Vereinigung bilden. Der Beweis zeigt, daß wir dadurch irgendwann fertig werden: T^* enthält ein Element f , zu dem kein f^* mehr existiert, und f oder f^{-1} ist dann die gesuchte Injektion.

Sehr viel direkter, transparenter und ebenso streng mathematisch kann der Algorithmus des Abtragens mit Hilfe der Ordinalzahlen durchgeführt werden, die wir im zweiten Abschnitt besprechen, und die in der Mengenlehre eine zentrale Rolle spielen. Die Ordinalzahlen erlauben uns, von den einzelnen Stufen des Abtragens zu sprechen: Es treten hierbei Nachfolgerstufen auf – die Stufe von f^* ist die Nachfolgerstufe der Stufe von f – und daneben Limesstufen, die der Vereinigung der Stufen von unendlich vielen f entsprechen.

Die Situation ist vergleichbar mit den beiden Beweisen des Äquivalenzsatzes von Bernstein und Zermelo. Der erste verläuft Schritt für Schritt, „von unten“, induktiv, der zweite über eine Schnittbildung „von oben“. Der obige strenge Beweis des Vergleichbarkeitssatzes ist nun ebenfalls ein Beweis „von oben“. Für eine rigorose Form des hier nur intuitiv geführten Beweises von unten haben wir an dieser Stelle das notwendige Werkzeug, die Ordinalzahlen – oder genauer: die transfiniten Zahlen –, nicht zur Verfügung.

Die Unterschiede der beiden Beweistypen „von unten“ und „von oben“ kann man beschreiben mit den beiden Möglichkeiten einen Koffer für eine Reise zu packen:

1. Man nimmt alles mit, was man braucht („von unten“).
2. Man läßt alles da, was man nicht braucht („von oben“).

Die zweite Möglichkeit ist legitim, aber doch etwas verschoben.

An dieser Stelle würde sich eine informelle Diskussion der Ordinalzahlen oder transfiniten Zahlen und des Zählens über die natürlichen Zahlen hinaus anbieten. Wir bleiben jedoch dem Thema *Größe* dieser Einführung treu, und besprechen das Thema *Ordnung* dann im zweiten Abschnitt.

Ein Satz von Felix Bernstein

Der Autor befürchtet, daß auch viele ehrgeizige Leser beim ersten Anlauf in der Mitte des Zermeloschen Beweises steckengeblieben sind. Dieser kurze Zwischenabschnitt bringt einen Satz von Felix Bernstein aus seiner Doktorarbeit von 1901 [siehe Bernstein 1905], der die Vergleichbarkeit zweier Mengen aus einer speziellen Voraussetzung gewinnt. Der transparente Beweis ist für sich hübsch und trickreich, und soll an dieser Stelle vor Augen führen, daß hinter dem komplizierten Beweis von Zermelo ein langer Weg zurückliegt, auf dem Teilresultate bereits einen Erfolg bedeuteten. Vielleicht ist das Motivation genug, sich dem Beweis von Zermelo anderntags noch einmal zu nähern.

Satz (*Trick von Felix Bernstein*)

Seien M, N Mengen, und es gelte $|M \times N| \leq |M \cup N|$.

Dann gilt $|M| \leq |N|$ oder $|N| \leq |M|$.

Beweis (*ohne Verwendung des Vergleichbarkeitssatzes von Zermelo*)

Sei $f : M \times N \rightarrow M \cup N$ injektiv.

1. Fall: Es existiert ein $x \in M$ mit $f''(\{x\} \times N) \subseteq M$.

Wir fixieren ein solches x und setzen $g(y) = f(x, y)$ für $y \in N$.

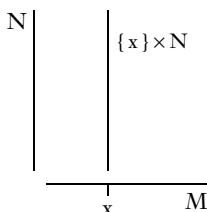
Dann ist $g : N \rightarrow M$ injektiv, also $|N| \leq |M|$.

2. Fall: Für alle $x \in M$ ist $B(x) = \{y \in N \mid f(x, y) \in N\} \neq \emptyset$.

Für $x \in M$ sei

$g(x) = f(x, y)$, wobei y ein beliebiges Element von $B(x) \subseteq N$ ist.

– Dann ist $g : M \rightarrow N$ injektiv, da f injektiv ist, und somit ist $|M| \leq |N|$.



Im ersten Fall sendet $f : M \times N \rightarrow M \cup N$ eine ganze Sektion von $M \times N$ der Form $\{x\} \times N$ in die Menge M . Aus $|N| = |\{x\} \times N|$ folgt dann $|N| \leq |M|$.

Andernfalls liegt auf jeder Sektion $\{x\} \times N$ mindestens ein Punkt, den f nach N schickt. Wir wählen dann aus jeder Sektion einen solchen Punkt y_x aus, und erhalten über den Weg: x nach y_x nach $f(x, y_x)$ eine Injektion von M nach N .

Der Beweis verwendet wieder eine Auswahl der Form „ein ...“, nämlich zur Definition der Funktion g .

Mit der an dieser Stelle merkwürdigen Voraussetzung $|M \times N| \leq |M \cup N|$ werden wir uns in Kapitel 12 beschäftigen. Schließlich sei bemerkt, daß das Argument des Beweises nicht nur historisch von Interesse ist, sondern auch in einem Satz von Alfred Tarski (1902 – 1983) 1924 wiederkehrt, der die Mächtigkeit von $M \times M$ mit Auswahlakten der Form „ein ...“ in Verbindung bringt (Satz von Bernstein-Tarski). Auch hierauf kommen wir noch zurück (2.5); für jetzt genügt uns der Satz von Bernstein als Dokument des Ringens um einen Beweis des Vergleichbarkeitssatzes. Daß in der Mathematik schöne Argumente selten bedeutungslos zu sein scheinen, ist darüber hinaus ein beruhigender Gedanke.

Zusammenfassung

Zwei wesentliche Resultate haben wir über das Konzept $|A| = |B|$ bzw. $|A| \leq |B|$ gezeigt: Den Äquivalenzsatz von Cantor und Bernstein und den Vergleichbarkeitssatz von Zermelo. Es fehlt uns noch der Nachweis der Reichhaltigkeit des Konzepts: Gibt es viele verschiedene Mächtigkeiten, d.h. eine Vielzahl von Mengen, die paarweise nicht gleichmächtig zueinander sind?

In unserer endlichen Realität haben wir viele verschieden große Nußhaufen – und ihre Mächtigkeiten entsprechen ideell den natürlichen Zahlen. Wie sieht es aber mit unendlichen Mengen aus?

Die Cantorsche Entdeckung, daß verschiedene Mächtigkeiten auch im Reich der unendlichen Mengen existieren, wurde zum Ausgangspunkt einer Entwicklung, die die Mathematik und das Bild der Mathematik tiefgreifend veränderte. Dem Feuerwerk zur neuen Zeit folgte zwar die obligatorische Katerstimmung (13. Kapitel), und ebenso zwangsläufig zogen daraufhin mathematische Wanderprediger mit neuen Dogmen und Verboten durchs Land; am Ende aber trat alles im neuen Glanz ans Licht, und heute bildet „der kühle Wurf“ der Mengenlehre das stattlichste Gebäude zur Diskussion von Grundlagenfragen der Mathematik. Unendlichkeit ist hier überall – außer im Eingangsbereich – ein gesellschaftsfähiges Thema.

Wir kommen im nächsten Kapitel zur Frage *Was ist unendlich?* und zum Vergleich der Mächtigkeiten verschiedener unendlicher Mengen. Dieses Kapitel, in dessen Zentrum der große Zermelosche Beweis steht, wäre jedoch ohne einen Blick auf Georg Cantor noch unvollständig.

Georg Cantor und das Vergleichbarkeitsproblem

Georg Cantor war von der Gültigkeit des Satzes immer überzeugt. In seiner Arbeit von 1878 heißt es bereits:

Cantor (1878): „Sind die beiden Mannigfaltigkeiten M und N nicht von gleicher Mächtigkeit, so wird entweder M mit einem Bestandteile von N oder es wird N mit einem Bestandteile von M die gleiche Mächtigkeit haben...“

Fast 30 Jahre hat es gedauert, bis der Vergleichbarkeitssatz streng bewiesen werden konnte. Cantor hat ein formales Argument für die Vergleichbarkeit nie veröffentlicht. In Briefen an Dedekind und Hilbert aus den Jahren 1897 – 1899 skizziert er jedoch eine Rechtfertigung des transfiniten Abtragealgorithmus, und mit dessen Hilfe einen Beweis des Vergleichbarkeitssatzes. Das Hauptproblem hierbei ist der Nachweis, daß man mit dem Abtragen einer Menge (oder zweier Mengen) „irgendwann fertig wird“. Cantor hat nun für diesen Nachweis genau jene Paradoxien der naiven Mengenlehre gewinnbringend eingesetzt, die wenige Jahre später von seinen Zeitgenossen als rein desaströs empfunden worden sind, und gerade von ihren Wiederentdeckern Cesare Burali-Forti (1861 – 1931) und Russell in einer Weise interpretiert wurden, daß die Interpretationen heute bedrohlicher erscheinen als die Paradoxien selber. Cantors Skizzen zu dem Problem sind dagegen Zeugen seiner überragend klaren Sicht der komplexen Phänomene, die der Mengenbegriff dem Wissensdurstigen bietet. Und wenn auch Cantors mathematische Kraft ab dem Zeitsprung ins 20. Jahrhundert zu technisch sauberen Zeichnungen dessen, was er sah, nicht mehr ausreichte, so haben seine Vorarbeiten dem ersten strengen Beweis des Vergleichbarkeitssatzes doch zumindest den Weg geebnet. Zermelo, der Ingenieur unter den Mengentheoretikern, konnte, die Skizzen Cantors in der Hand, den entscheidenden Schritt weiter gehen, und mit Hilfe der Theorie der Wohlordnungen den Satz im Jahre

1904 beweisen. Vier Jahre später fand er dann einen Weg, das Getriebe der Wohlorderungen aus seinem Beweis zu entfernen, und der in seiner Arbeit von 1908 verwendeten Argumentation folgt der obige komplexe, aber begrifflich reduzierte Beweis des Vergleichbarkeitssatzes.

Die Geschichte des Satzes ist ein Beispiel für die Nichtlinearität des mathematischen Fortschritts: Erst eine recht weitgehende Entwicklung der Theorie machte es möglich, eine Frage elementar zu beantworten, die gleich zu Beginn der Untersuchung des Mächtigkeitsbegriffs auftritt. Wir können heute das Problem unmittelbar nach seiner bloßen Formulierbarkeit lösen, mit dem ererbten Wissen derer, die nach einem mühevollen Aufstieg die Küstenlinien schließlich klar erkennen konnten.

Georg Cantor über das Problem der Vergleichbarkeit

„Wir haben gesehen, daß [für zwei Kardinalzahlen α , β] von den drei Beziehungen

$$\alpha = \beta, \alpha < \beta, \beta < \alpha$$

jede einzelne die beiden anderen ausschließt.

Dagegen versteht es sich keineswegs von selbst und dürfte an dieser Stelle unseres Gedankenganges kaum zu beweisen sein, daß bei irgend zwei Kardinalzahlen α und β eine von jenen drei Beziehungen notwendig realisiert sein müsse.

Erst später, wenn wir einen Überblick über die aufsteigende Folge der transfiniten Kardinalzahlen und eine Einsicht in ihren Zusammenhang gewonnen haben werden, wird sich die Wahrheit des Satzes ergeben:

A. „Sind α und β zwei beliebige Kardinalzahlen, so ist entweder $\alpha = \beta$ oder $\alpha < \beta$ oder $\alpha > \beta$.“

(Georg Cantor 1895, „Beiträge zur Begründung der transfiniten Mengenlehre. I“)

6. Unendliche Mengen

Wir haben $|A| = |B|$ als „A und B haben die gleiche Größe“ interpretiert, wobei „gleiche Größe“ durch die Existenz einer Bijektion zwischen A und B definiert wird. Gilt $|A| = |B|$, so ist eine vollständige Paarbildung der Elemente von A und B möglich, die Elemente der Mengen entsprechen sich vollkommen, wenn wir eine bestimmte Abbildung zugrunde legen.

Es taucht nun das merkwürdige und zunächst irritierende Phänomen auf, daß manche Mengen gleich groß sind zu einem echten Teil von sich selbst: Es gibt Mengen A, B mit $|A| = |B|$ und $B \subset A$.

Man kann dieses Phänomen zur Definition der Unendlichkeit einer Menge benutzen, die nur von den elementaren Mächtigkeitsbegriffen und nicht von einer irgendwie definierten Anzahl der Elemente einer Menge Gebrauch macht, und wir werden diesem Weg folgen. Vor einer derartigen Definition stellen wir zur Motivation eine an dieser Stelle fast schon etwas naive Frage, und betrachten dann noch ein merkwürdiges Gebäude, das sogenannte „Hilbertsche Hotel“.

Die frühe Literatur zur Mengenlehre verwendete viele didaktisch ambitionierte Seiten auf die Diskussion der vermeintlich paradoxen Eigenschaften unendlicher Mengen, und vor Cantor schrieb Bolzano ein ganzes Buch über „Paradoxien des Unendlichen“. Heute sind wir durch eine viel strenger aufgebaute und logisch durchtrainierte Mathematik daran gewöhnt, daß Begriffe wie „ist größer als“ oder „unendlich“ erst durch Definition zu Begriffen der Mathematik werden. Die Intuition spielt für die Formulierung einer Definition eine große Rolle, hat sich dann aber an deren Konsequenzen zu orientieren und nicht umgekehrt. Die sorgfältige mathematische Entfaltung einer Definition zeigt zum Glück zumeist, daß hier keine subtile Gehirnwäsche stattfindet, sondern eine Aufklärung im besten platonischen Sinne: Schattenhafte verschwommene Bilder werden schließlich farbig und scharf umrissen. Die Erfahrung ist die des Findens und Entdeckens und nicht die des eigentlichen Erschaffens oder auch nur des freien Gestaltens nach festen Spielregeln.

Gibt es mehr natürliche oder mehr gerade Zahlen ?

Seien $\mathbb{N} = \{0, 1, 2, 3, \dots\}$ und $\mathbb{G} = \{0, 2, 4, 6, \dots\}$ die Menge der geraden Zahlen. Definiere $f : \mathbb{N} \rightarrow \mathbb{G}$ durch

$$f(n) = 2n \text{ für } n \in \mathbb{N}.$$

Dann ist $f : \mathbb{N} \rightarrow \mathbb{G}$ bijektiv, also gilt $|\mathbb{N}| = |\mathbb{G}|$.

Offenbar ist aber $\mathbb{G} \subset \mathbb{N}$.

Ist das Ergebnis *$\mathbb{G}$ und $\mathbb{N}$ sind gleich groß* nicht paradox, wo doch wegen $\mathbb{G} \subset \mathbb{N}$ die Menge $\mathbb{G}$ offensichtlich kleiner ist als $\mathbb{N}$? Keineswegs, denn hier liegen ver-

schiedene Interpretationen von „groß“ vor. Beide sind natürlich, aber sie stimmen im allgemeinen nicht überein.

A ist größergleich als B falls *B ist Teilmenge von A* ist ein sinnvoller Größenbegriff, und er wird in der Mengenlehre oft verwendet. Er ist aber vom Begriff der Größe, der durch Bijektionen gegeben wird, verschieden, und hinsichtlich des Zieles, daß die Größe einer Menge Zahlcharakter haben soll, ist er unbrauchbar: Die Vergleichbarkeit $A \subseteq B$ oder $B \subseteq A$ gilt nicht für beliebige Mengen A und B.

Hausdorff (1914): „Unsere Beispiele [...] zeigten ja, daß eine Menge recht wohl mit einer ihrer echten Teilmengen äquivalent sein kann, z. B. die Menge der natürlichen Zahlen n mit der Menge der Quadratzahlen n^2 . Diese Eigenschaft kann offenbar nur unendlichen Mengen zukommen und kommt ihnen, wie wir sehen werden [...], auch stets zu. Wenn wir also den Zeichen $= < >$ die übliche Bedeutung lassen und insbesondere verlangen wollen, daß von den drei Fällen

$$\alpha = b, \quad \alpha < b, \quad \alpha > b$$

immer nur einer eintreten kann, so werden wir darauf verzichten müssen, jeder echten Teilmenge von A eine Kardinalzahl $< \alpha$ zu geben; wir müssen das geheiligte Axiom ‚totum parte majus‘ verletzen, wie wir uns überhaupt darauf gefaßt machen müssen, daß die Rechnung mit unendlichen Kardinalzahlen in vielen Punkten von der mit endlichen abweichen wird, ohne daß darin der geringste Einwand gegen diese unendlichen Zahlen zu erblicken wäre.“

Das Hilbertsche Hotel

Fast schon zur mathematische Folklore geworden ist das Hilbertsche Hotel. Dieses Hotel hat für jede natürliche Zahl ein Zimmer:



Alle Zimmer sind belegt. Ein neuer Gast kann aber wie folgt untergebracht werden:

- (i) Jeder alte Gast zieht von Zimmer n nach Zimmer $n + 1$.
- (ii) Der neue Gast wird in Zimmer 0 einquartiert.

Derartige Möglichkeiten des Platzmachens durch Verschiebung sind gerade charakteristisch für unendliche Mengen. Es ist vielleicht ein Vergnügen für den Leser sich zu überlegen, wie er neue Gäste $G_0, G_1, \dots, G_n, \dots, n \in \mathbb{N}$, die alle gleichzeitig ankommen, in einem bereits ausgebuchten Hilbertschen Hotel unterbringen würde.

Dedekinds Definition der Unendlichkeit

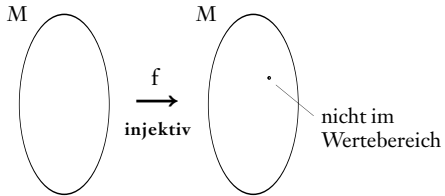
Definition (Unendlichkeit nach Dedekind)

Sei M eine Menge. M heißt *unendlich*, falls es eine echte Teilmenge N von M gibt, die sich bijektiv auf M abbilden läßt, d. h. es gibt ein $N \subset M$ mit $|N| = |M|$.

Eine Menge heißt *endlich*, falls sie nicht unendlich ist.

Anders formuliert: Eine Menge M ist unendlich, falls es eine Injektion $f : M \rightarrow M$ gibt, die mindestens einen Wert ausläßt, d. h. $\text{rng}(f) \neq M$.

Nach obigem Beispiel ist $\mathbb{N}$ eine unendliche Menge – wie es sein soll. Allgemeiner gilt:



Übung

Sei $A \subseteq \mathbb{N}$ unbeschränkt in $\mathbb{N}$, d. h. für alle $n \in \mathbb{N}$ gibt es ein $m \in A$ mit $n \leq m$. Dann ist A unendlich.

[Wie definieren $g : A \rightarrow A$ durch $g(n) = \text{„das kleinste } m \in A \text{ mit } n < m\text{“}$.]

Dedekind (1888): „Ein System [Menge] S heißt unendlich, wenn es einem echten Teile seiner selbst ähnlich [gleichmächtig] ist; im entgegengesetzten Falle heißt S ein endliches System ...

[Fußnote zu dieser Definition:] ... In dieser Form habe ich die Definition des Unendlichen, welche den Kern meiner ganzen Untersuchung bildet, im September 1882 Herrn G. Cantor und schon mehrere Jahre früher auch den Herren Schwarz und Weber mitgeteilt. Alle anderen mir bekannten Versuche, das Unendliche vom Endlichen zu unterscheiden, scheinen mir so wenig gelungen zu sein, daß ich auf eine Kritik derselben verzichten zu dürfen glaube.“

Der Leser vergleiche hierzu auch Bolzanos „Paradoxien des Unendlichen“ (1851), §21f.

Wir werden unten eine Reihe von „gelungenen“ äquivalenten Definitionen von *unendlich* und *endlich* kennenlernen, wobei die ansprechendste unter ihnen die natürlichen Zahlen verwendet. Dedekinds Definition bleibt in ihrem Purismus ungeschlagen, auch wenn sich im axiomatischen Aufbau der Mengenlehre eine Definition über die natürlichen Zahlen als einfacher erweist, im Sinne des geringeren Gewichts der die Definition inhaltlich tragenden Axiome.

Cantor hat, obwohl er mit dem Phänomen hinter der Dedekindschen Definition und mit ihr selbst vertraut war, viel kompliziertere, aber dafür auch sehr interessante Merkmale der Endlichkeit an die Spitze gestellt. So schreibt er in seiner philosophischen Arbeit von 1887:

Cantor (1887): „Unter einer *endlichen* [*nichtleeren*] Menge verstehen wir eine solche M , welche aus *einem* ursprünglichen Element durch sukzessive Hinzufügung neuer Elemente derartig hervorgeht, daß auch *rückwärts* aus M durch *sukzessive* Entfernung der Elemente *in umgekehrter Ordnung* das ursprüngliche Element gewonnen werden kann...

Als *durchaus wesentliches* Merkmal *endlicher* Mengen muß es angesehen werden, daß eine solche *keinem ihrer* [*echten*] Bestandteile äquivalent ist. Denn eine aktual unendliche Menge ist *immer* so beschaffen, daß auf mehrfache Weise ein Bestandteil von ihr bezeichnet werden kann, der ihr äquivalent ist.“

Cantors endliche Mengen sind also den Stapeln der Informatik ähnlich, die durch schrittweises „push“ an Höhe gewinnen, aber auch durch schrittweises „pop“ wieder reduziert werden können. Ein Stapel unendlicher Höhe bestehend aus allen natürlichen Zahlen $n \in \mathbb{N}$ in ihrer natürlichen Ordnung ist ideell vorstellbar. Man kann ihn sich durch sukzessives „push“ aufgebaut denken, dagegen kann er durch „pop“ nicht mehr von oben abgebaut werden, weil er kein oberstes Objekt mehr besitzt.

In der Vorstellung Cantors sind Mengen zwar extensional bestimmt, aber dennoch oft mit einer inneren Ordnung versehen. Konzeptionell wie intuitiv sind heute Mengen nackt und ungeordnet. Auch Ordnungen auf ihnen, wie etwa $<_{\mathbb{N}}$, sind, für sich selbst genommen, ungeordnete Mengen.

Dedekind hat aus seiner Definition die „Unendlichkeit der Gedankenwelt“ abgeleitet (vgl. das Zitat auf dem Vorblatt des Buches):

Hessenberg (1906): „Einer der interessantesten Versuche, die Existenz transfiniter Mengen zu beweisen, ist der von Dedekind unternommene. Es sei a irgend ein Gegenstand des Denkens, so kann ich das Urteil fällen: a ist ein Gegenstand meines Denkens. Dieses Urteil $\varphi(a)$ ist selbst ein Gegenstand des Denkens. Die Zuordnung φ zwischen a und $\varphi(a)$ ist umkehrbar eindeutig [injektiv] und bildet die Menge aller Gedankendinge auf einen echten Teil ihrer selbst ab, da nicht jeder Gegenstand des Denkens die Form eines Urteils, daher a fortiori nicht die Form des speziellen Urteils $\varphi(a)$ hat. Demnach ist die Menge aller Gedankendinge transfinit.“

Die Idee ist hier gerade, die natürlichen Zahlen nicht zu verwenden. Sonst könnte man analog argumentieren: Die Zuordnung φ , die $n \in \mathbb{N}$ auf $n + 1$ abbildet, ist injektiv. Natürliche Zahlen kommen erst später, die Unendlichkeit wohnt dem Denken selber inne, und dieses braucht nicht erst das Zählen, um sich dieser Tatsache bewußt zu werden. Mathematisch läßt sich das Argument sicher nicht als Beweis der Existenz unendlicher Mengen auffassen. Philosophisch ist die Idee zweifellos interessant, und kulturgeschichtlich ist der Versuch, die Existenz transfiniter Mengen aus einem iterierten Akt der Selbstreflexion zu beweisen, ein schönsten Beispiel aufklärerischen Denkens. Der Mensch erkennt durch bloße Selbstbeobachtung die unendlichen Möglichkeiten seines Verstandes. Nicht zufällig ist es Hessenberg, der Dedekinds Versuch würdigt: Hessenberg war philosophisch gesehen Neukantianer. (Sein Buch von 1906 erschien zuerst in den „Abhandlungen der Friesschen Schule“.)

Einfache Sätze über unendliche Mengen

Wir leiten aus der Dedekind-Definition einige elementare Resultate ab.

Satz (*Übertragung der Unendlichkeit zwischen Mengen gleicher Mächtigkeit*)

Seien A und B gleichmächtige Mengen. Dann gilt:

Ist A unendlich, so ist auch B unendlich.

Beweis

Sei $A' \subset A$, und sei $f : A \rightarrow A'$ bijektiv. Weiter sei $h : A \rightarrow B$ bijektiv.

Wir setzen

$$g = h \circ f \circ h^{-1} : B \rightarrow B$$

Dann ist g injektiv. Ist $x \in A - A'$, so ist $h(x) \notin \text{rng}(g)$.

(Genauer gilt $\text{rng}(g) = h''A' \subset h''A = B$.)

– Also ist $g : B \rightarrow \text{rng}(g) \subset B$ ein Zeuge für die Unendlichkeit von B .

Der nächste Satz zeigt die Übertragung der Unendlichkeit auf jede Obermenge einer unendlichen Menge:

Satz (*Übertragung der Unendlichkeit auf Obermengen*)

Seien A, B Mengen, und es gelte $A \subseteq B$. Dann gilt:

Ist A unendlich, so ist auch B unendlich.

Beweis

Sei $A' \subset A$ und $f : A \rightarrow A'$ bijektiv. Sei

$$g = f \cup \text{id}_{B-A}.$$

(Also $g(b) = f(b)$ für $b \in A$, $g(b) = b$ für $b \in B - A$.)

Dann ist g injektiv.

Ist $x \in A - A'$, so ist $x \notin \text{rng}(g)$.

(Genauer gilt $\text{rng}(g) = A' \cup (B - A) \subset B$.)

– Also ist $g : B \rightarrow \text{rng}(g) \subset B$ ein Zeuge für die Unendlichkeit von B .

Als Korollar zu diesen beiden Sätzen erhalten wir:

Korollar (*Übertragung der Unendlichkeit auf gleichmächtige und größere Mengen*)

Seien A, B Mengen, und es gelte $|A| \leq |B|$. Dann gilt:

Ist A unendlich, so ist auch B unendlich.

Beweis

Sei $h : A \rightarrow B$ injektiv, und sei $C = \text{rng}(h)$.

Dann ist $|A| = |C|$, also ist C unendlich.

– Aber $C \subseteq B$, also ist auch B unendlich.

Interessanter sind die Reduktionen von unendlichen Mengen, die die Unendlichkeit erhalten. Zunächst zeigen wir, daß ein Tropfen an der Unendlichkeit des Meeres nichts ändert.

Satz (*Entfernen eines Elementes*)

Sei A eine unendliche Menge. Weiter seien $a \in A$ und $B = A - \{a\}$.
Dann ist auch B unendlich.

Beweis

Sei $A' \subset A$ und $f: A \rightarrow A'$ bijektiv.

Sei $b \in A - A' (\neq \emptyset)$. Wir setzen:

$$g = f|_{(A - \{b\})}.$$

Dann ist g injektiv, und $f(b) \notin \text{rng}(g)$.

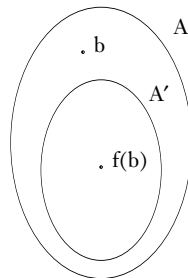
Aber $f(b) \neq b$ wegen $\text{rng}(f) = A'$. Also ist

$$g: A - \{b\} \rightarrow A - \{b, f(b)\} \subset A - \{b\}$$

ein Zeuge für die Unendlichkeit von $A - \{b\}$.

Aber offenbar $|A - \{a\}| = |A - \{b\}|$, also ist auch

– $A - \{a\}$ unendlich nach dem Satz oben.



Korollar (*Hinzufügen eines Elementes*)

Seien B eine endliche Menge und a ein beliebiges Objekt.

Weiter sei $A = B \cup \{a\}$.

Dann ist auch A endlich.

Beweis

Andernfalls ist A unendlich (und $a \notin B$). Nach dem Satz oben ist

– dann $A - \{a\} = B$ unendlich, *im Widerspruch* zur Voraussetzung.

Wiederholte Anwendung des Korollars ergibt, daß $B \cup \{a_1, \dots, a_n\}$ endlich ist für endliche B und für beliebige Objekte $a_1, \dots, a_n, n \in \mathbb{N}$. Insbesondere sind also (für $B = \emptyset$) alle Mengen der Form $\{a_1, \dots, a_n\}$ endlich. Wir wissen noch nicht, daß umgekehrt alle endlichen Mengen die Gestalt $\{a_1, \dots, a_n\}, n \in \mathbb{N}$, haben; wir werden dies unten zeigen.

Für weitergehende Resultate wird die rein funktionale Definition der Unendlichkeit im rein funktionalen Kontext recht schwerfällig. Der Nachweis, daß die Vereinigung zweier – oder stärker endlich vieler – endlicher Mengen wieder endlich ist, bereitet bereits Schwierigkeiten. (Der Leser mag versuchen, dies im Stil der obigen Beweise zu zeigen). An dieser Stelle kommen uns nun die natürlichen Zahlen zu Hilfe, ähnlich wie in der Mächtigkeitstheorie zum Beweis des Satzes von Cantor-Bernstein. Wie dort wäre es eher künstlich, die Stärke rekursiver Definitionen und induktiver Beweise beim Heben schwererer Gewichte nicht zu benutzen.

Unendlichkeit und natürliche Zahlen

Zeugen für die Unendlichkeit einer Menge A sind Injektionen von A nach A , die Werte auslassen. Andererseits ist eine solche Injektion auf ganz A definiert, insbesondere also auch auf den Werten, die sie selbst ausläßt. Dies führt zur folgenden allgemeinen Definition.

Definition (Orbit eines Punktes)

Sei $g : A \rightarrow A$ eine Funktion, und sei $x \in A$.

Dann ist $S_x : \mathbb{N} \rightarrow A$, der *Orbit* von x unter g in A , rekursiv wie folgt definiert:

$$S_x(0) = x,$$

$$S_x(n+1) = g(S_x(n)) \quad \text{für } n \in \mathbb{N}.$$

Ein Orbit S heißt *azyklisch*, falls S injektiv ist. Andernfalls heißt S *zyklisch*.

Ein $x \in A$ heißt *azyklisch*, falls S_x azyklisch ist. Andernfalls heißt x *zyklisch*.

Der Buchstabe S erinnert hierbei an „Spur“. Der Orbit S_x von x unter g beschreibt die Bahn des Punktes x , wenn wiederholt die Funktion g auf x und seine Bilder angewendet wird.

Ist $g(x) = x$, so ist $S_x(n) = x$ für alle $n \in \mathbb{N}$. Ist $g(x) = y$, $g(y) = x$ und $x \neq y$, so ist $S_x(n) = x$ für gerade n und $S_x(n) = y$ für ungerade n . Das Wort „zyklisch“ wird zudem motiviert durch die folgende Übung.

Übung (Orbitalbahn eines zyklischen Punktes)

Sei $g : A \rightarrow A$ eine Funktion, und sei $x \in A$ zyklisch. Dann gilt:

Es existieren $n_0 \in \mathbb{N}$ und $m \in \mathbb{N} - \{0\}$ mit der Eigenschaft:

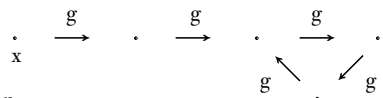
Für alle $n, n' \geq n_0$ gilt:

$$S_x(n) = S_x(n') \quad \text{gdw} \quad n \equiv_m n'.$$

Ist g injektiv, so ist $n_0 = 0$ geeignet.

[Zur Erinnerung: $n \equiv_m n'$ gdw n und n' haben den gleichen Rest bei Division durch m .]

Das (eindeutig bestimmte) derartige m heißt dann der *Zyklus* von x , das kleinste derartige n_0 die *Vorlaufzeit* von x .



Eine weitere, etwas informal formulierte Übungsaufgabe für den Leser ist, sich die möglichen Orbit-Typen unter nicht injektiven und unter surjektiven Funktionen $g : A \rightarrow A$ zu überlegen. Der „Weg rückwärts“ von x zu einem y mit $g(y) = x$ ist für surjektive g immer möglich, aber im allgemeinen nicht eindeutig. Für Injektionen dagegen kann man weiter den Rückwärtsorbit eines Punktes x definieren (solange entsprechende Urbilder existieren). Für Bijektionen gibt es dann stets einen unendlichen Vorwärts- und Rückwärtsorbit, und die Orbitalbahn wird am besten durch ein $h : \mathbb{Z} \rightarrow A$ beschrieben mit $h(0) = x$. Diese $\mathbb{Z}$ -Orbits für Bijektionen sind dann entweder Kreise oder $\mathbb{Z}$ -ähnlich ohne Überlappung. Der Leser ist aufgefordert, auf dem Papier ein bißchen herumzuspielen. Er kann hier ein Stück Mathematik selber formulieren und entdecken.

Der Begriff des Orbits ist für sich interessant, spielt aber im Umfeld der Unendlichkeit nach Dedekind eine besondere Rolle:

Satz (*Existenz azyklischer Orbits*)

Sei $g : A \rightarrow A$ injektiv.

Dann ist jedes $x \in A - \text{rng}(g)$ azyklisch.

Beweis

Sei $x \in A - \text{rng}(g)$, und sei $S = S_x$ der Orbit von x unter g .

Wir zeigen durch Induktion nach n :

Für alle $n \in \mathbb{N}$ gilt: $S(n) \neq S(m)$ für alle $0 \leq m < n$.

Induktionsschritt n : Annahme $S(n) = S(m)$ für ein $0 \leq m < n$.

Dann ist $m \neq 0$ wegen $S(0) = x \notin \text{rng}(g)$ und $S(n) \in \text{rng}(g)$.

Dann aber $S(n-1) = S(m-1)$ wegen g injektiv,

– *im Widerspruch* zur Induktionsvoraussetzung.

Den Leser hat der Begriff des Orbits und der Beweis des obigen Satzes vielleicht an den Beweis des Satzes von Cantor-Bernstein erinnert. In der Tat haben wir dort bereits implizit die Orbits aller Punkte der Menge C unter der Funktion f betrachtet (vgl. das Diagramm im Beweis des Satzes).

Ist also A eine unendliche Menge, so gibt es nach dem Existenzsatz eine Injektion $g : A \rightarrow A$ und ein $x \in A$, dessen Orbit unter g azyklisch ist. Die Umkehrung ist Teil des folgenden Satzes, der den fundamentalen Zusammenhang zwischen unendlichen Mengen und natürlichen Zahlen wiedergibt.

Satz (*Einbettbarkeit der natürlichen Zahlen in unendliche Mengen*)

Sei A eine Menge. Dann sind äquivalent:

- (i) A ist unendlich.
- (ii) Es gibt ein $g : A \rightarrow A$ und ein $x \in A$ mit einem azyklischen Orbit.
- (iii) $|\mathbb{N}| \leq |A|$.

Beweis

(i) $\hookrightarrow$ (ii): Sei $g : A \rightarrow A' \subset A$ injektiv.

Dann ist jedes $x \in A - A'$ azyklisch unter g nach dem Satz oben.

(ii) $\hookrightarrow$ (iii): Sei S ein azyklischer Orbit unter g . Dann ist $S : \mathbb{N} \rightarrow A$ injektiv.

– (iii) $\hookrightarrow$ (i): $\mathbb{N}$ ist unendlich. Wegen $|\mathbb{N}| \leq |A|$ also auch A .

Die letzte Implikation kann man auch so beweisen:

Sei $h : \mathbb{N} \rightarrow A$ injektiv. Für $n \in \mathbb{N}$ sei $x_n = h(n)$.

Wir „verschieben die x_n um eins“: Definiere $g : A \rightarrow A$ durch

$$g(x) = \begin{cases} x_{n+1}, & \text{falls } x = x_n \text{ für ein } n \in \mathbb{N}, \\ x & \text{sonst.} \end{cases}$$

Dann ist $g : A \rightarrow A - \{x_0\}$ bijektiv. Also A unendlich.

Weiter gilt: h ist der Orbit von x_0 unter g .

Mit diesen Ergebnissen können wir nun zeigen, daß wir unendliche Mengen nicht durch endliche Vereinigungen von endlichen Mengen erhalten können.

Satz (*Endliche Zerlegungen unendlicher Mengen*)

Sei A eine unendliche Menge, und sei $P \subseteq \mathcal{P}(A)$ eine endliche Zerlegung von A , d.h. P ist endlich, $\bigcup P = A$, und die Elemente von P sind paarweise disjunkte Mengen. Dann existiert ein unendliches $B \in P$.

Beweis (*durch Besuchszeitenanalyse*)

Sei $S : \mathbb{N} \rightarrow A$ der Orbit eines azyklischen $x \in A$ (unter einer beliebigen Injektion $g : A \rightarrow A' \subset A$). Für $B \in P$ sei

$$N_B = \{n \in \mathbb{N} \mid S(n) \in B\}.$$

[N_B ist die Menge der „Besuchszeiten“ von x in der Menge B .]

Ist N_B unbeschränkt in $\mathbb{N}$ für ein $B \in P$, so ist $|\mathbb{N}| = |N_B| \leq |B|$, also B unendlich, und wir sind fertig.

Annahme, N_B ist beschränkt in $\mathbb{N}$ für alle $B \in P$.

Für $B \in P$ mit $N_B \neq \emptyset$ sei

$$m_B = \max(N_B) = \text{„das größte } m \in \mathbb{N} \text{ mit } m \in N_B\text{“}.$$

[Die Menge B wird also zum Zeitpunkt m_B zum letzten Mal besucht.]

Schließlich sei

$$U = \{m_B \mid B \in P, N_B \neq \emptyset\}.$$

Dann ist U unbeschränkt in $\mathbb{N}$.

[*Annahme*, es gibt ein $k \in \mathbb{N}$ mit $m < k$ für alle $m \in U$.

Sei $B \in P$ mit $S(k) \in B$. Dann ist $k \in N_B$, also $k \leq m_B \in U$, *Widerspruch!*]

Aber $f : U \rightarrow P$ mit

$$f(m) = \text{„das eindeutige } B \in P \text{ mit } S(m) \in B\text{“} \quad \text{für } m \in U$$

ist injektiv nach Konstruktion von U .

– Also $|\mathbb{N}| = |U| \leq |P|$. Also ist P unendlich, *Widerspruch!*

Übung

Seien A, B endliche Mengen. Dann gilt:

- (i) $A \cup B$ ist endlich,
- (ii) $A \times B$ ist endlich.

Satz (*Endliche Überdeckungen unendlicher Mengen*)

Sei A eine unendliche Menge, und sei $P \subseteq \mathcal{P}(A)$ eine endliche Überdeckung von A , d.h. P ist endlich und $\bigcup P = A$. Dann existiert ein unendliches $B \in P$.

Beweis

Fast genau wie eben. Lediglich die Definition der Funktion f lautet nun:

– $f(m) = \text{„ein } B \in P \text{ mit } S(m) \in B \text{ und } m_B = m\text{“}$ für $m \in U$.

Korollar (*Vereinigung endlich vieler endlicher Mengen*)

Sei P eine endliche Menge, und jedes $B \in P$ sei endlich.

Dann ist $\bigcup P$ endlich.

In der Analyse des Unendlichkeitsbegriffs nach Dedekind haben wir eine Auswahl der Form „ein ...“ zum ersten Mal für den Überdeckungssatz oben verwendet. In der Tat ist ein „ein ...“ hier (und für einen Beweis des Korollars) unvermeidlich. Daß die Vereinigung endlich vieler paarweise disjunkter Mengen endlich ist, läßt sich ohne Auswahlakte beweisen, da dies ein Korollar zum Zerlegungssatz ist. Weiter gilt $B_1 \cup B_2 = B_1 \cup (B_2 - B_1)$, und es folgt induktiv ganz ohne „ein ...“, daß $B_1 \cup \dots \cup B_n$ endlich ist für beliebige endliche Mengen B_i , $1 \leq i \leq n$, $n \in \mathbb{N}$.

Dagegen ist für die folgende Übung wieder eine Definition der Form „ein ...“ nötig:

Übung

Sei A eine endliche Menge. Dann ist $\mathcal{P}(A)$ endlich.

[Sei $S: \mathbb{N} \rightarrow \mathcal{P}(A)$ ein azyklischer Orbit. Jedes $S(n) \subseteq A$ ist endlich.

Für alle $\{a_1, \dots, a_k\} \subseteq A$, $k \in \mathbb{N}$, existiert ein $n \in \mathbb{N}$ mit

$$S(n) - \{a_1, \dots, a_k\} \neq \emptyset,$$

denn S ist injektiv und $\mathcal{P}(\{a_1, \dots, a_k\})$ hat nur 2^k Elemente.

Eine induktive Auswahl aus solchen nichtleeren Mengen

liefert einen Widerspruch zur Endlichkeit von A .]

Endlichkeit und natürliche Zahlen

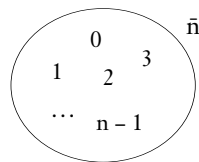
Das offene Problem aus dem letzten Zwischenabschnitt lautete: Ist jede endliche Menge von der Form $\{a_1, \dots, a_n\}$ für ein $n \in \mathbb{N}$? Dies wollen wir nun beweisen.

Zunächst definieren wir für jedes $n \in \mathbb{N}$ eine Referenzmenge mit genau n Elementen.

Definition (*die Mengen $\bar{n}$ für $n \in \mathbb{N}$*)

Für $n \in \mathbb{N}$ sei

$$\bar{n} = \{0, 1, \dots, n-1\} = \{m \in \mathbb{N} \mid m < n\}.$$



Die in der Realität auftauchenden Nußhaufen sind endlich im Sinne der Definition von Dedekind. Andernfalls wäre der Algorithmus des paarweisen Abtragens der Nußhaufen H_1 und H_2 keine korrekte Methode zur Feststellung eines Größenunterschiedes zwischen den Haufen!

Die Mengen $\bar{n}$ sind nun gewissermaßen die Nußhaufen der mathematischen Welt, und für sie können wir die Korrektheit des Abtragealgorithmus beweisen (vgl. das Korollar unten).

Satz (Endlichkeit von $\bar{n}$)

Für jedes $n \in \mathbb{N}$ ist $\bar{n}$ endlich.

Beweis

Wir zeigen die Aussage durch Induktion nach $n \in \mathbb{N}$.

Induktionsanfang $n = 0$:

$\bar{0} = \{ \}$ ist offenbar endlich.

Induktionsschritt von n nach $n + 1$:

Sei $m = n + 1$. Dann ist $\bar{m} = \{ 0, \dots, n \} = \bar{n} \cup \{ n \}$.

Nach Induktionsvoraussetzung ist $\bar{n}$ endlich. Nach dem Satz

- über das Hinzufügen eines Elementes ist dann auch $\bar{m}$ endlich.

Wir geben noch einen direkten, von den vorangehenden Überlegungen unabhängigen Beweis des Satzes.

Zweiter Beweis des Satzes

Annahme nicht. Sei dann n das kleinste Element von $\mathbb{N}$ mit „ $\bar{n}$ ist unendlich“.

Sei dann $f : \bar{n} \rightarrow A \subset \bar{n}$ bijektiv.

Offenbar ist $n \neq 0$ und $A \neq \emptyset$. Sei also $n = n' + 1$.

O.E. ist $n' \in A$.

Andernfalls sei $A' = (A - \{ \min(A) \}) \cup \{ n' \}$.

Dann gilt $A' \subset \bar{n}$ und $|A| = |A'|$. Also existiert

$g : \bar{n} \rightarrow A'$ bijektiv, $A' \subset \bar{n}$, und es ist $n' \in A'$.

O.E. ist zudem $f(n') = n'$.

Andernfalls existiert wegen $n' \in A$ ein $n^* < n'$ mit $f(n^*) = n'$.

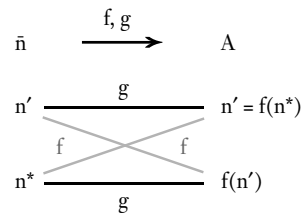
Wir setzen dann

$$g = (f - \{ (n', f(n')), (n^*, n') \}) \cup \{ (n', n'), (n^*, f(n')) \}.$$

g ist genau wie f , lediglich die beiden

Werte für n' und n^* sind vertauscht!

Dann ist $g : \bar{n} \rightarrow A$ bijektiv und $g(n') = n'$.



Wegen $f : \bar{n} \rightarrow A$ bijektiv und $f(n') = n'$ ist dann aber

$f|_{\bar{n}'} : \bar{n}' \rightarrow A - \{ n' \}$ bijektiv, und $A - \{ n' \} \subset \bar{n}'$.

- Also ist $\bar{n}'$ unendlich, *im Widerspruch* zur minimalen Wahl von n .

Korollar ($|\bar{n}| < |\bar{m}|$ für $n < m$)

Seien $n, m \in \mathbb{N}$, und sei $f : \bar{n} \rightarrow \bar{m}$ bijektiv. Dann gilt $n = m$.

Beweis

Sei $f : \bar{n} \rightarrow \bar{m}$ bijektiv, $n \neq m$. O.E. ist $m < n$ (sonst betrachte f^{-1}).

- Dann ist aber $\bar{m} \subset \bar{n}$, also $\bar{n}$ unendlich, *Widerspruch!*

Übung (Taubenschlagprinzip)

Sei $n \in \mathbb{N}$, und sei $f: \bar{n} \rightarrow \bar{n}$. Dann sind äquivalent:

- (i) f ist injektiv.
- (ii) f ist bijektiv.
- (iii) f ist surjektiv.

[Ist f surjektiv, so betrachte $g: \bar{n} \rightarrow \bar{n}$ mit $g(m) =$ „das kleinste k mit $f(k) = m$ “.]

Schließlich erhalten wir eine nützliche Äquivalenz für die Endlichkeit einer Menge:

Satz (Charakterisierung der endlichen Mengen mit Hilfe der natürlichen Zahlen)

Sei A eine Menge. Dann sind äquivalent:

- (i) A ist endlich.
- (ii) Es existiert ein $n \in \mathbb{N}$ mit $|A| = |\bar{n}|$.

Beweis

(i) $\hookrightarrow$ (ii): Sei A endlich. Ist $A = \emptyset$ sind wir fertig. Andernfalls sei $x_0 \in A$.

Wir definieren rekursiv für $n \in \mathbb{N}$:

$x_{n+1} =$ „ein $x \in A - \{x_0, \dots, x_n\}$ “,

solange $A - \{x_0, \dots, x_n\} \neq \emptyset$.

[Hier begegnet uns wieder eine Auswahl-Situation.]

Die erhaltenen x_n sind paarweise verschieden. Also muß ein kleinstes $n^* \in \mathbb{N}$ existieren, für das x_{n^*} nicht definiert ist.

[Denn *andernfalls* ist $g: \mathbb{N} \rightarrow A$ mit $g(n) = x_n$ injektiv, also $|\mathbb{N}| \leq |A|$, im Widerspruch zu A endlich.]

Dann ist aber $f: \bar{n}^* \rightarrow A$ mit $f(n) = x_n$ bijektiv. Also gilt (ii).

(ii) $\hookrightarrow$ (i): Sei $n \in \mathbb{N}$ und $f: A \rightarrow \bar{n}$ bijektiv.

Wäre A unendlich, so wäre auch $\bar{n}$ unendlich,

— im Widerspruch zum Satz oben. Also ist A endlich.

Es folgt, daß jede endliche Menge A die Form $A = \{a_1, \dots, a_n\}$ für ein $n \in \mathbb{N}$ hat, denn es gilt $|A| = |\bar{n}|$ für ein $n \in \mathbb{N}$, und dann ist $A = \{f(0), \dots, f(n-1)\}$ für eine beliebige Bijektion $f: \bar{n} \rightarrow A$.

Der Beweis des Satzes verwendet die Sätze über Zerlegungen und Überdeckungen unendlicher Mengen nicht, und wir erhalten neue Beweise für die Endlichkeit der Vereinigung endlicher vieler endlicher Mengen. Diese Beweise sind sehr einfach, ruhen aber auf dem obigen Auswahlakt im Beweis von (i) $\hookrightarrow$ (ii).

Übung

Zeigen Sie mit Hilfe des obigen Charakterisierungssatzes:

- (i) Die Vereinigung zweier endlicher Mengen ist endlich.
- (ii) Die Vereinigung endlich vieler endlicher Mengen ist endlich.

Definition ($\mathbb{N}$ -unendlich)

Eine Menge A heißt $\mathbb{N}$ -endlich, falls es ein $n \in \mathbb{N}$ gibt mit $|A| = |\bar{n}|$.
 A heißt $\mathbb{N}$ -unendlich, falls A nicht $\mathbb{N}$ -endlich ist.

Man kann die Theorie umgekehrt aufziehen, und die $\mathbb{N}$ -Versionen als primäre Definitionen der Unendlichkeit und Endlichkeit benutzen, was ebenso natürlich erscheint wie ein Start mit den Dedekind-Definitionen. Wir haben gezeigt, daß eine Menge genau dann Dedekind-endlich ist, wenn sie $\mathbb{N}$ -endlich ist. Interessant ist, daß ein Beweis der Implikation „Dedekind-endlich *folgt* $\mathbb{N}$ -endlich“ oder gleichwertig „ $\mathbb{N}$ -unendlich *folgt* Dedekind-unendlich“ abstrakte Auswahlakte benötigt. Die Unterschiede der beiden Definitionen werden sehr klar, wenn man folgende Umformulierung von $\mathbb{N}$ -unendlich betrachtet:

Übung

Sei A eine Menge. Dann sind äquivalent:

- (i) A ist $\mathbb{N}$ -unendlich.
- (ii) $|\bar{n}| \leq |A|$ für alle $n \in \mathbb{N}$.

Dedekind-unendlich ist dagegen gleichwertig mit $|\mathbb{N}| \leq |A|$. Die kritische Implikation lautet dann also umformuliert:

Ist $|\bar{n}| \leq |A|$ für alle $n \in \mathbb{N}$, so gilt $|\mathbb{N}| \leq |A|$.

Das sieht sehr überzeugend aus, braucht aber bereits starke Hilfsmittel für einen Beweis. Mit ähnlichen Feinheiten wollen wir uns nun beschäftigen.

Andere Charakterisierungen der Unendlichkeit

Es gibt eine ganze Reihe weiterer Äquivalenzen zur Endlichkeit und Unendlichkeit, und dieser Zwischenabschnitt bringt hierzu vier Beispiele, ohne einen Sport aus dem Thema zu machen. Der Leser, der mit der obigen Diskussion bereits zufrieden ist, kann ihn überspringen, da die hier diskutierten Begriffe später nicht mehr verwendet werden. Die Existenz vieler verschiedener Definitionen, die sich dann als gleichwertig herausstellen, ist ein sicheres Indiz dafür, daß ein besonders natürliches Konzept zur Diskussion steht. (Der Leser vergleiche die vielen äquivalenten Definitionen des Begriffs „berechenbar/rekursiv“ in den Lehrbüchern zur mathematischen Logik.)

Obwohl alle Definitionen, die wir geben werden, sich als äquivalent herausstellen, taucht hier in vielen Fällen das gleiche Phänomen auf wie oben: Die Hälfte der Äquivalenz ist elementar (wenn auch zuweilen trickreich), die andere benötigt Auswahlakte. Die folgenden Definitionen sind nun so geordnet, daß die Implikation von einer Definition auf die folgende (also „ $\wedge$ “) elementar ist, die Rückrichtung (also „ $\vee$ “) aber abstrakte Auswahlakte der Form „ein ...“ benötigt. Eine Ausnahme bildet die letzte Definition nach Alfred Tarski, die in beiden Richtungen elementar äquivalent zur $\mathbb{N}$ -Unendlichkeit ist. Wir starten mit Dedekind-unendlich, und gelangen schließlich über mehrere Schritte zur

$\mathbb{N}$ - bzw. Tarski-Unendlichkeit. Es ist interessant, daß sich zwischen den intuitiv schon sehr nahe beieinanderliegenden Aussagen „ $|\mathbb{N}| \leq |A|$ “ und „ $|\bar{n}| \leq |A|$ “ für alle $n \in \mathbb{N}$ mehrere natürliche Zwischenstufen finden lassen.

Im folgenden definieren wir immer nur *XYZ-endlich* oder *XYZ-unendlich*. Durch Verneinung der Aussagen erhält man dann *XYZ-unendlich* bzw. *XYZ-endlich*.

Beim Betrachten der Dedekind-Definition bietet sich der folgende Begriff an:

Definition (*Dedekind*-unendlich*)

Sei A eine Menge. A heißt *Dedekind*-unendlich*, falls ein surjektives $f : A \rightarrow A$ existiert, das nicht injektiv ist.

Dedekind*-unendlich ist also gerade invers zur Dedekind-Unendlichkeit: Die Bedingung an $f : A \rightarrow A$ lautet *injektiv-nichtsurjektiv* für Dedekind-unendlich, und *surjektiv-nichtinjektiv* für Dedekind*-unendlich.

Übung

Sei A eine Menge. Dann sind äquivalent:

- (i) A ist Dedekind-unendlich.
- (ii) A ist Dedekind*-unendlich.

Man kann weiter die Charakterisierung $|\mathbb{N}| \leq |A|$ der Dedekind-Unendlichkeit von A invertieren:

Definition (*Dedekind**-unendlich*)

Sei A eine Menge. A heißt *Dedekind**-unendlich*, falls ein surjektives $f : A \rightarrow \mathbb{N}$ existiert.

Satz

Sei A eine Menge. Dann sind äquivalent:

- (i) A ist Dedekind*-unendlich.
- (ii) A ist Dedekind**-unendlich.

Beweis

(i) $\hookrightarrow$ (ii): Sei $g : A \rightarrow A$ surjektiv und nicht injektiv.

Seien $a, b, c \in A$ mit $b \neq c$ und $g(b) = g(c) = a$.

Für $x \in A$ sei wieder S_x der Orbit von x unter g .

Es gilt $b \notin \text{rng}(S_a)$ oder $c \notin \text{rng}(S_a)$ (!).

O. E. sei $b \notin \text{rng}(S_a)$. Wir definieren dann $f : A \rightarrow \mathbb{N}$ durch

$$f(x) = \begin{cases} \text{„das kleinste } n \text{ mit } S_x(n) = b\text{“}, & \text{falls ein solches } n \text{ existiert,} \\ 0 & \text{sonst.} \end{cases}$$

Annahme $\text{rng}(f) \neq \mathbb{N}$. Sei dann n minimal mit $n \notin \text{rng}(f)$.

Wegen $f(b) = 0$ ist $n > 0$. Nach minimaler Wahl existiert ein x mit

$f(x) = n - 1$. Dann ist $S_x(n - 1) = b$.

Wegen g surjektiv existiert ein y mit $g(y) = x$.

Dann ist $S_y(n) = S_x(n-1) = b$, und somit $0 \leq f(y) \leq n$.

Nach Annahme ist $f(y) \neq n$.

Aber auch $f(y) \neq 0$, da sonst $y = b$, $x = a$ und $b = S_a(n-1) \in \text{rng}(S_a)$ wäre.

Also $0 < f(y) < n$. Dann ist $S_x(f(y)-1) = S_y(f(y)) = b$.

Also $f(x) \leq f(y) - 1 < n - 1$, *Widerspruch*.

(ii) $\cap$ (i): Sei $f: A \rightarrow \mathbb{N}$ surjektiv. Für $n \in \mathbb{N}$ sei

$x_n =$ „ein $x \in A$ mit $f(x) = n$ “.

Sei $X = \{x_n \mid n \geq 1\}$. Definiere $g: A \rightarrow A$ durch

$g(x_n) = x_{n-1}$ für $n \in \mathbb{N} - \{0\}$, $g(x) = x$ für $x \notin X$.

— Dann ist $g: A \rightarrow A$ surjektiv, aber $g(x_0) = g(x_1)$, $x_0 \neq x_1$.

Die Funktion f in (i) $\cap$ (ii) nimmt immer größere Werte an, je weiter wir von b entlang Urbildern zurückgehen, was im allgemeinen wegen mangelnder Injektivität ein netzartiges Bild ergibt. $b \notin \text{rng}(S_a)$, $g(b) = a$ liefert, daß wir beim Zurückgehen nie mehr auf b selbst stoßen, und die Existenz eines solchen b ist genau die Stelle im Beweis, wo wir mehr brauchen als die Surjektivität von g .

Einen schnellen Beweis des Satzes erhält man natürlich, wenn man die Äquivalenz von $|\mathbb{N}| \leq |A|$ und $|\mathbb{N}| \leq^* |A|$, d.h. „es gibt ein $f: A \rightarrow \mathbb{N}$ surjektiv“, benutzt: Dann ist Dedekind**-unendlich äquivalent mit Dedekind-unendlich, und nach der Übung oben dann auch mit Dedekind*-unendlich. Auf diese Weise haben wir (i) $\cap$ (ii) über die Kette *Dedekind*-unendlich* $\cap$ *Dedekind-unendlich* $\cap$ *Dedekind**-unendlich* gezeigt, für die erste Implikation aber Auswahlakte verwendet; der obige Beweis von (i) $\cap$ (ii) ist davon frei.

Eine weitere Definition der Unendlichkeit stammt von Alfred Tarski (1924). Sie betrifft die $\subseteq$ -Relation auf der Potenzmenge $\mathcal{P}(A)$ einer Menge A : Erlaubt ein Teilmengensystem $P \subseteq \mathcal{P}(A)$ unbegrenzt viele Schritte, die von einem $x \in P$ zu einem $y \in P$ führen mit $x \subset y$, so treffen wir bei jedem Schritt neue Elemente von A an, und damit ist A dann intuitiv unendlich. Wir definieren hierzu:

Definition (*letztes Glied einer Kette*)

Sei K eine $\subseteq$ -Kette. K hat ein *letztes Glied*, falls es ein $x \in K$ gibt mit:

$y \subseteq x$ für alle $y \in K$.

Definition (*Tarski-Unendlichkeit und Ketten-Unendlichkeit*)

(a) Eine Menge A heißt *Ketten-endlich*, falls gilt:

Jede nichtleere Kette $K \subseteq \mathcal{P}(A)$ hat ein letztes Glied.

(b) Eine Menge A heißt *Tarski-endlich*, falls gilt:

Für jedes nichtleere $P \subseteq \mathcal{P}(A)$ gibt es ein $y \in P$ mit:

$\text{non}(y \subset x)$ für alle $x \in P$.

Eine äquivalente und sprachlich kompaktere Formulierung der Bedingung in (b) ist: Jedes nichtleere $P \subseteq \mathcal{P}(A)$ besitzt ein $\subseteq$ -maximales Element.

Übung

Jedes unbeschränkte $A \subseteq \mathbb{N}$ ist Ketten-unendlich.

Übung

Seien A, B gleichmächtige Mengen. Dann gilt:

- (i) Ist A Ketten-endlich, so ist B Ketten-endlich.
- (ii) Ist A Tarski-endlich, so ist B Tarski-endlich.

Satz

Sei A eine Menge. Dann sind äquivalent:

- (i) A ist Dedekind**-unendlich.
- (ii) A ist Ketten-unendlich.

Beweis

(i) $\hookrightarrow$ (ii):

Sei $f: A \rightarrow \mathbb{N}$ surjektiv. Für $n \in \mathbb{N}$ sei

$$A_n = f^{-1}n = \{x \in A \mid f(x) < n\}.$$

Dann gilt $A_n \subset A_m$ für alle $n < m$.

Also ist $\{A_n \mid n \in \mathbb{N}\} \subseteq \mathcal{P}(A)$ eine Kette ohne letztes Glied.

(ii) $\hookrightarrow$ (i):

Sei $K \subseteq \mathcal{P}(A)$ eine nichtleere Kette ohne letztes Glied.

Sei $X_0 \in K$ beliebig. Wir definieren $X_n \in K$ für $n \geq 1$ rekursiv durch

$$X_{n+1} = \text{„ein } X \in K \text{ mit } X_n \subset X\text{“}.$$

Weiter definieren wir $f: A \rightarrow \mathbb{N}$ durch

$$f(x) = \begin{cases} \text{„das kleinste } n \text{ mit } x \in X_n\text{“,} & \text{falls ein solches } n \text{ existiert,} \\ 0, & \text{sonst.} \end{cases}$$

– Dann ist $f: A \rightarrow \mathbb{N}$ surjektiv.

Die beiden noch fehlenden Äquivalenzen behandeln die folgenden Übungen:

Übung

Sei A eine Menge. Dann sind äquivalent:

- (i) A ist Ketten-unendlich.
- (ii) A ist Tarski-unendlich.

[Die nichttriviale Richtung (ii) $\hookrightarrow$ (i) verwendet eine Auswahldefinition zur Gewinnung einer nichtleeren Kette ohne ein letztes Glied aus einem Zeugen $P \subseteq \mathcal{P}(A)$ für die Tarski-Unendlichkeit von A .]

Übung

Sei A eine Menge. Dann sind äquivalent:

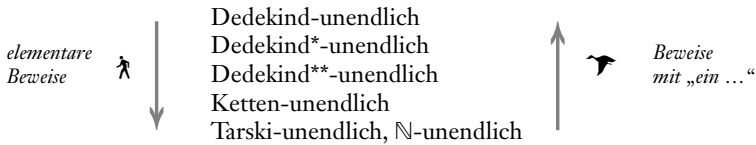
- (i) A ist Tarski-unendlich.
- (ii) A ist $\mathbb{N}$ -unendlich.

[Für die Richtung (i) $\hookrightarrow$ (ii) zeige „ $\bar{n}$ ist Tarski-endlich“ durch Induktion.

Für (ii) $\hookrightarrow$ (i) betrachte $P = \{X \subseteq A \mid |X| = |\bar{n}| \text{ für ein } n \in \mathbb{N}\}$.

Beide Richtungen verwenden keine Auswahlakte.]

Im Sinne der ewigen Wiederkunft des Auswahlthemas erhalten wir damit nun das folgende Bild:



Der Autor hofft, daß der vorangehende logische Abstieg von *Dedekind-unendlich* zu *$\mathbb{N}$ -unendlich* dem Leser eine abwechslungsreiche Wanderung gewesen ist. Abwärts geht es zu Fuß, aufwärts brauchen wir einen Sessellift, mit Ausnahme der elementar äquivalenten unendlichen Talstationen *$\mathbb{N}$ -unendlich* und *Tarski-unendlich*.

Um die tatsächliche Notwendigkeit der Auswahlakte zu beweisen, braucht man viel weitergehende Techniken; es könnte ja ein einfacher Beweis übersehen worden sein.

Wir schließen dieses Kapitel mit einem Auszug aus einem Vortrag von David Hilbert, gehalten am 4. Juni 1925 in Münster anläßlich der „Ehrung des Andenkens an Weierstraß“ (Karl Weierstraß 1815 – 1897).

David Hilbert über Unendlichkeit

„... Durch diese Bemerkungen wollte ich nur dartun, daß die endgültige Aufklärung über das *Wesen des Unendlichen* weit über den Bereich spezieller fachwissenschaftlicher Interessen vielmehr zur *Ehre des menschlichen Verstandes* selbst notwendig geworden ist.

Das Unendliche hat wie keine andere Frage von jeher so tief das *Gemüt* der Menschen bewegt; das Unendliche hat wie kaum eine andere *Idee* auf den Verstand so anregend und fruchtbar gewirkt; das Unendliche ist aber auch wie kein anderer *Begriff* so der *Aufklärung* bedürftig.

Wenn wir uns nun dieser Aufgabe, das Wesen des Unendlichen aufzuklären, zuwenden, so müssen wir uns in aller Kürze vergegenwärtigen, welche inhaltliche Bedeutung dem Unendlichen in der Wirklichkeit zukommt; wir sehen zunächst, was wir aus der Physik darüber erfahren.

Der erste naive Eindruck von dem Naturgeschehen und der Materie ist der des stetigen, des Kontinuierlichen. Haben wir ein Stück Metall oder ein Flüssigkeitsvolumen, so drängt sich uns die Vorstellung auf, daß sie unbegrenzt teilbar seien, daß ein noch so kleines Stück von ihnen immer wieder dieselben Eigenschaften habe. Aber überall, wo man die Methoden der Forschung in der Physik der Materie genügend verfeinerte, stieß man auf Grenzen für die Teilbarkeit, die nicht an der Unzulänglichkeit unserer Versuche, sondern in der Natur der Sache liegen, so daß man geradezu die Tendenz der modernen Wissenschaften als eine Emanzipierung von dem Unendlichkleinen auffassen könnte und daß man jetzt an Stelle des alten Leitsatzes: ‚*natura non facit saltus*‘ das Gegenteil ‚die Natur macht Sprünge‘ behaupten könnte.

Bekanntlich ist alle Materie aus kleinen Bausteinen, den Atomen zusammengesetzt ... Während bis dahin die Elektrizität als ein Fluidum galt ..., so erwies sich jetzt auch sie aufgebaut aus positiven und negativen *Elektronen* ... Nun selbst die Energie läßt, wie heute feststeht, die unendliche Zerteilung nicht schlechthin und uneingeschränkt zu; Planck entdeckte die *Energiequanten*.

Und das Fazit ist jedenfalls, daß ein homogenes Kontinuum, welches die fortgesetzte Teilbarkeit zuließe, und somit das Unendliche im Kleinen realisieren würde, in der Wirklichkeit nirgends angetroffen wird. Die unendliche Teilbarkeit eines Kontinuums ist nur eine in Gedanken vorhandene Operation, nur eine Idee, die durch unsere Beobachtungen der Natur und die Erfahrungen der Physik und Chemie widerlegt wird.

Die zweite Stelle, an der uns in der Natur die Frage nach der Unendlichkeit entgegentritt, treffen wir bei der Betrachtung der Welt als Ganzes. Hier haben wir die Ausdehnung der Welt zu untersuchen, ob es in ihr ein Unendlichgroßes gibt.

Die Meinung von der Unendlichkeit der Welt war lange Zeit die herrschende; bis zu Kant und auch weiterhin noch hegte man an der Unendlichkeit des Raumes überhaupt keinen Zweifel.

Hier ist es wieder die moderne Wissenschaft, insbesondere die Astronomie, die diese Frage von neuem aufrollt und sie nicht durch das unzulängliche Hilfsmittel metaphysischer Spekulation, sondern durch Gründe, die sich auf die Erfahrung stützen und auf der Anwendung von Naturgesetzen beruhen, zu entscheiden sucht. Und es haben sich schwerwiegende Einwände gegen die Unendlichkeit herausgestellt. Zur Annahme der Unendlichkeit des Raumes führt mit Notwendigkeit die *Euklidische* Geometrie. Nun ist zwar die Euklidische Geometrie ein in sich widerspruchsfreies Gebäude und Begriffssystem; daraus folgt aber noch nicht, daß sie in der Wirklichkeit Gültigkeit besitzt. Ob dies der Fall ist, kann allein die Beobachtung und Erfahrung entscheiden ... Einstein hat die Notwendigkeit gezeigt, von der Euklidischen Geometrie abzugehen. Auf Grund seiner Gravitationstheorie nimmt er auch die kosmologischen Fragen in Angriff und zeigt die Möglichkeit einer endlichen Welt, und alle von den Astronomen gefundenen Resultate sind auch mit [dieser] Annahme ... verträglich.

Die Endlichkeit des Wirklichen haben wir nun in zwei Richtungen festgestellt: nach dem Unendlichkleinen und dem Unendlichgroßen. Dennoch könnte es sehr wohl zutreffen, daß das Unendliche *in unserem Denken* einen wohlberechtigten Platz hat und die Rolle eines unentbehrlichen Begriffes einnimmt ... “

(David Hilbert, 1925. Auch in: David Hilbert 1964, „Hilbertiana“)

7. Abzählbare Mengen

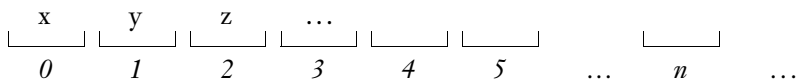
Die einfachste unendliche Menge, die wir kennen, ist die Menge $\mathbb{N}$ der natürlichen Zahlen. Die natürlichen Zahlen werden insbesondere zum Zählen, Indizieren, Durchnumerieren, Ordnen, etc. anderer Mengen verwendet. Dies führt zu folgendem Begriff, der alle Mengen beschreibt, die sich durch natürliche Zahlen in dieser Weise kontrollieren lassen:

Definition (*abzählbare Mengen*)

Eine Menge M heißt *abzählbar*, falls gilt:

- (i) es existiert ein bijektives $f: \bar{n} \rightarrow M$ für ein $n \in \mathbb{N}$, oder
- (ii) es existiert ein bijektives $f: \mathbb{N} \rightarrow M$.

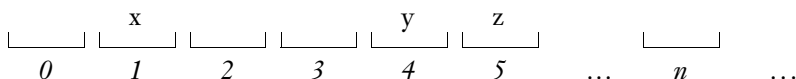
Der Begriff „abzählbar“ entspricht folgender Anschauung. Eine Menge M ist abzählbar, wenn wir alle Elemente $x, y, z, \dots$ von M der Reihe nach auf $\mathbb{N}$ -viele Plätze verteilen können:



Dabei ist für unendliche Mengen eine „geschickte“ Platzanweisung notwendig: Ist etwa $M = \mathbb{N}$, und besetzen wir Platz n mit der Zahl $2n$, so bleiben „am Ende“ die ungeraden Zahlen stehen, obwohl $\mathbb{N}$ sicher abzählbar ist. Aus den Resultaten des vorangehenden Kapitels folgt, daß für endliche Mengen ein geschickter Platzanweiser entbehrlich ist: Nach jeder Platzzuweisung einer Menge mit genau n Elementen ist immer der Platz n der erste freie Sitz.

Die $\mathbb{N}$ -Endlichkeit, die wir zur Definition von abzählbar verwendet haben, scheint der intuitiven Bedeutung des Wortes „abzählen“ näher zu sein als die (äquivalente) Dedekind-Endlichkeit.

Wir können auch Platzanweisungen betrachten, bei denen einige Plätze freibleiben dürfen, bevor ein neuer Platz besetzt wird:



Eine solche lückenhafte Verteilung kann bequem durch eine Injektion von einer Menge in die natürlichen Zahlen beschrieben werden. Durch ein intuitiv klares, aber recht aufwendig zu organisierendes Nachrückverfahren – man denke an

ein dunkles Kino – erhalten wir „richtige“ Abzählungen. Einfacher ist es an den Sitzen entlangzugehen und dabei M neu abzuzählen, wie dies im Beweis des folgenden Satzes geschieht.

Satz (*Abzählbarkeit und Einbettbarkeit in die natürlichen Zahlen*)

Sei M eine Menge. Dann sind äquivalent:

- (i) M ist abzählbar.
- (ii) $|M| \leq |\mathbb{N}|$.

Beweis

(i) $\hookrightarrow$ (ii):

Ist M abzählbar, so ist $|M| = |\bar{n}|$ für ein $n \in \mathbb{N}$ oder $|M| = |\mathbb{N}|$.
In beiden Fällen ist offenbar $|M| \leq |\mathbb{N}|$.

(ii) $\hookrightarrow$ (i):

Sei $|M| \leq |\mathbb{N}|$ und $f: M \rightarrow \mathbb{N}$ injektiv.

Wir zählen nun $\text{rng}(f) \subseteq \mathbb{N}$ monoton auf. Hierzu definieren wir durch Rekursion über $n \in \mathbb{N}$ solange möglich eine Funktion g wie folgt:

$g(n) =$ „das kleinste $k \in \text{rng}(f)$ mit $k > g(i)$ für alle $0 \leq i < n$ “,
falls ein solches k existiert.

Sei $N = \text{dom}(g)$. Dann ist $g: N \rightarrow \text{rng}(f)$ bijektiv.

Also ist $f^{-1} \circ g: N \rightarrow M$ bijektiv.

Aber es gilt $N = \bar{n}$ für ein $n \in \mathbb{N}$ oder $N = \mathbb{N}$ (!).

– Also ist M abzählbar.

Zu einem schnellen Beweis führt natürlich eine Unterscheidung „ $\text{rng}(f)$ endlich“, also $|M| = |\text{rng}(f)| = |\bar{n}|$ für ein $n \in \mathbb{N}$, oder „ $\text{rng}(f)$ unendlich“ also $|M| = |\text{rng}(f)| = |\mathbb{N}|$ zu einer Funktion g wie gewünscht, die dann aber mit f i. a. nichts mehr zu tun hat. Unsere Funktion g erhält die durch f gegebene Sitzordnung. Ein derartiges strukturelles Zusammenschieben eines Wertebereichs ist in der Mengenlehre vielfach nützlich.

Es gilt $f^{-1} \circ g = h^{-1}$ mit $h(x) =$ „das $n \in \mathbb{N}$ mit $|\bar{n}| = |\{y \in M \mid f(y) < f(x)\}|$ “ für $x \in M$.

Wir zeigen nun, daß jede unendliche Menge eine abzählbar unendliche Teilmenge besitzt. „Abzählbar unendlich“ ist also die erste „Anzahl“ nach den endlichen Größen.

Satz

Sei M eine unendliche Menge.

Dann existiert eine abzählbar unendliche Teilmenge von M .

Beweis

Wegen M unendlich gilt $|\mathbb{N}| \leq |M|$ nach dem Satz über die Einbettbarkeit der natürlichen Zahlen in unendliche Mengen.

Sei also $f: \mathbb{N} \rightarrow M$ injektiv.

– Dann ist $\text{rng}(f)$ eine abzählbar unendliche Teilmenge von M .

Der Beweis des Satzes ist elementar, da M als (Dedekind)-unendlich vorausgesetzt wird. Jedoch ist für den Beweis, daß jede $\mathbb{N}$ -unendliche Menge eine abzählbar unendliche Teilmenge enthält, ein Auswahlakt notwendig. Man definiert hierzu rekursiv:

$x_n =$ „ein $x \in M$ mit $x \neq x_i$ für alle $0 \leq i < n$ “.

Aus M $\mathbb{N}$ -unendlich folgt, daß diese Rekursion nicht abbricht. $X = \{x_n \mid n \in \mathbb{N}\}$ ist dann eine abzählbar unendliche Teilmenge von M .

Die entscheidende Frage ist nun:

Ist jede Menge abzählbar?

Anders formuliert:

Ist jede unendliche Menge M gleichmächtig zu den natürlichen Zahlen, d. h. existiert für jede unendliche Menge M ein bijektives $f: \mathbb{N} \rightarrow M$?

Das Konzept der Mächtigkeit wäre dann nicht besonders interessant, denn jede Menge wäre dann entweder vom „Größentyp“ $\bar{n} = \{0, \dots, n-1\}$ für ein $n \in \mathbb{N}$, oder aber vom Typ $\mathbb{N}$.

Es zeigt sich zunächst: Viele prominente Mengen sind abzählbar. Dies wollen wir jetzt von den ganzen, den rationalen und den algebraischen Zahlen zeigen. Zudem führen viele Operationen mit abzählbaren Mengen nicht aus dem Reich der Abzählbarkeit heraus.

$\mathbb{Z}$ ist abzählbar

Wir betrachten die Menge $\mathbb{Z} = \{\dots, -1, 0, 1, \dots\}$ der ganzen Zahlen und definieren $f: \mathbb{Z} \rightarrow \mathbb{N}$ durch

$$f(x) = \begin{cases} 2x, & \text{falls } x \geq 0, \\ -2x-1, & \text{falls } x < 0. \end{cases}$$

Dann ist $f: \mathbb{Z} \rightarrow \mathbb{N}$ bijektiv. Also ist $\mathbb{Z}$ abzählbar.

Allgemeiner gilt der folgende Satz:

Satz

Seien A und B abzählbar. Dann ist auch $A \cup B$ abzählbar.

Beweis

Seien $f: A \rightarrow \mathbb{N}$, $g: B \rightarrow \mathbb{N}$ injektiv. Definiere $h: A \cup B \rightarrow \mathbb{Z}$ durch

$$h(x) = \begin{cases} f(x), & \text{falls } x \in A, \\ -(g(x) + 1), & \text{falls } x \in B - A. \end{cases}$$

Dann ist h injektiv, also $|A \cup B| \leq |\mathbb{Z}| = |\mathbb{N}|$.

– Also ist $A \cup B$ abzählbar.

$\mathbb{N} \times \mathbb{N}$ ist abzählbar: Die Cantorsche Paarungsfunktion

Satz

$\mathbb{N} \times \mathbb{N}$ ist abzählbar.

Beweis

Definiere $\pi : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ bijektiv durch

$$\pi(a, b) = 1/2(a + b)(a + b + 1) + a.$$

Z. B. $\pi(3, 2) = 30/2 + 3 = 18.$

π heißt die *Cantorsche Paarungsfunktion*. π ist bijektiv und zählt die Menge $\mathbb{N} \times \mathbb{N}$ diagonal auf, wie in dem folgenden Diagramm dargestellt:

7								
6								
5	15							
4	10	16						
3	6	11	17					
2	3	7	12	18				
1	1	4	8	13	19			
0	0	2	5	9	14	20		
	0	1	2	3	4	5	6	7

Obige Formel für die Werte von π erhält man durch die Beobachtung, daß in der ersten Senkrechten des Diagramms ($a = 0$) die Partialsummen der Reihe $0 + 1 + 2 + 3 + 4 \dots$ stehen, und diese berechnen sich zu $1/2(n(n + 1))$. Weiter ist die Summe $a + b$ der Koordinaten konstant auf den Diagonalen.

Gegeben (a, b) bestimmt man zuerst die Diagonale, in der sich (a, b) befindet. Der erste Wert dieser Diagonalen ist nach obiger Überlegung $\pi(0, a + b) = 1/2(a + b)(a + b + 1)$. Und offenbar ist dann $\pi(a, b) = \pi(0, a + b) + a$.

Übung

Für alle $n \in \mathbb{N}$ gilt: $1 + 2 + \dots + n = (n(n + 1))/2$.

[Durch Induktion oder durch den Gaußtrick $n + 1 = (n - 1) + 2 = (n - 2) + 3 = \dots$]

Übung

- (i) Sei A abzählbar. Dann ist auch $A \times A$ abzählbar.
- (ii) Ist A abzählbar und $n \in \mathbb{N}$, $n \geq 1$, so ist auch A^n abzählbar.
- (iii) Gilt $|A \times A| = |A|$ für eine Menge A , so auch $|A^n| = |A|$ für alle $n \in \mathbb{N}$, $n \geq 1$.

Die Cantorsche Paarungsfunktion ist ein Polynom in a und b zweiten Grades. Es ist erstaunlich, daß die diagonale Aufzählung von $\mathbb{N} \times \mathbb{N}$ durch so eine einfache Funktion beschrieben werden kann.

Cantor (1878): „Es hat nämlich die Funktion $\mu + ((\mu + v - 1)(\mu + v - 2))/2$, wie leicht zu zeigen, die bemerkenswerte Eigenschaft, daß sie alle positiven ganzen Zahlen und jede nur einmal darstellt, wenn in ihr μ und v unabhängig voneinander ebenfalls jeden positiven, ganzzahligen Wert erhalten.“

In der Tat gibt es keine weiteren *Polynome höchstens zweiten Grades in a und b* , die $\mathbb{N} \times \mathbb{N}$ bijektiv auf $\mathbb{N}$ abbilden, mit Ausnahme der Funktion $\tilde{\pi} : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$, definiert durch $\tilde{\pi}(a, b) = \pi(b, a)$. Diese Tatsache, die die herausragende Stellung der Cantorschen Paarungsfunktion belegt, ist nichttrivial und als Satz von Fueter-Polya (1923) bekannt – zum Beweis wird das Transzendenz-Resultat von Ferdinand Lindemann (1852 – 1939) benutzt. Viele Fragen in diesem Umfeld sind noch offen. Z. B. ist unbekannt, ob es ein Polynom höheren Grades in a, b geben kann, das $\mathbb{N} \times \mathbb{N}$ bijektiv auf $\mathbb{N}$ abbildet [hierzu und zum Satz von Fueter-Polya siehe Smoryński 1991 und Lew/Rosenberg 1978].

Im Hinblick auf den Satz von Fueter-Polya ist interessant, daß es „Fast-Polynome“ zweiten Grades in a, b gibt, die $\mathbb{N} \times \mathbb{N}$ bijektiv auf $\mathbb{N}$ abbilden:

Übung

Definiere $\rho : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ durch

$$\rho(a, b) = \max(a^2, b^2) + a + (a \dot{-} b) \quad \text{für } a, b \in \mathbb{N}, \text{ wobei}$$

$$a \dot{-} b = \begin{cases} a - b, & \text{falls } b \leq a, \\ 0, & \text{falls } b > a. \end{cases}$$

(Man wird ρ lieb gewinnen, wenn man einige Werte in ein $\mathbb{N} \times \mathbb{N}$ -Gitter wie im Diagramm zur Cantorschen Paarungsfunktion einzeichnet.)

Es gilt: $\rho : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ ist bijektiv.

Die Konstruktion der Cantorschen Paarungsfunktion läßt sich auf höhere Dimensionen verallgemeinern, und man erhält ein bijektives $\pi_n : \mathbb{N}^n \rightarrow \mathbb{N}$ für alle $n \geq 3$. π_n ist ein Polynom n -ten Grades in $a_1, \dots, a_n$.

Eine zweite Möglichkeit, Bijektionen zwischen $\mathbb{N}^n$ und $\mathbb{N}$ zu erhalten, ist die Komposition, z. B.

$$\pi^3(a_1, a_2, a_3) = \pi(\pi(a_1, a_2), a_3),$$

$$\pi^4(a_1, a_2, a_3, a_4) = \pi(\pi^3(a_1, a_2, a_3), a_4),$$

$$\tilde{\pi}^4(a_1, a_2, a_3, a_4) = \pi(\pi(a_1, a_2), \pi(a_3, a_4)), \text{ usw.}$$

Beide Möglichkeiten liefern Polynome in $a_1, a_2, \dots, a_n$, die zweite erzeugt Polynome höheren Grades als die Anzahl der Variablen, z. B. ist π^3 vom Grad 4. Es ist nicht bekannt, ob es Polynome gibt, die $\mathbb{N}^n$ bijektiv auf $\mathbb{N}$ abbilden und die nicht durch diese beiden Möglichkeiten (und Umordnung der Variablen wie in $\tilde{\pi}(a, b) = \pi(b, a)$) gebildet werden.

Abzählungen von $\mathbb{N} \times \mathbb{N}$ und allgemeiner $\mathbb{N}^n$ für $n \geq 2$ gibt es natürlich zuhauf. Jede Wanderung auf dem $\mathbb{N} \times \mathbb{N}$ -Feld, in beliebigem Zickzackkurs, die jeden Punkt $(a, b) \in \mathbb{N} \times \mathbb{N}$ genau einmal besucht, liefert eine Bijektion zwischen $\mathbb{N}$ und $\mathbb{N} \times \mathbb{N}$.

Übung

Geben Sie ein Polynom dritten Grades in a, b, c an, welches $\mathbb{N}^3$ bijektiv auf $\mathbb{N}$ abbildet.

[Analog zur Cantorsche Aufzählung, wobei nun die $\mathbb{N}^3$ -Punkte der Ebenen $a + b + c = 0, a + b + c = 1, \dots$ nacheinander geeignet aufgezählt werden.]

Übung

$f: \mathbb{N}^2 \rightarrow \mathbb{N}$ mit $f(n, m) = 2^{n+1}(m+1) - 2^n - 1$ für alle $n, m \in \mathbb{N}$ ist bijektiv.

Primzahlen und $|\mathbb{N} \times \mathbb{N}| = |\mathbb{N}|$

Primzahlen liefern Zerlegungen von $\mathbb{N}$ in unendlich viele unendliche Teile, und hieraus erhält man weitere Beweise von $|\mathbb{N} \times \mathbb{N}| = |\mathbb{N}|$. Für $n \in \mathbb{N}$ sei

$p_n =$ „die $(n+1)$ -te Primzahl“,

also $p_0 = 2, p_1 = 3, p_2 = 5, p_3 = 7, p_4 = 11, \dots$

Dann sind für $n \in \mathbb{N}$ die Mengen $P_n = \{p_n^{k+1} \mid k \in \mathbb{N}\}$ paarweise disjunkte unendliche Teilmengen von $\mathbb{N}$ [Eindeutigkeit der Primfaktorzerlegung!]. Setzen wir

$f(n, k) = p_n^{k+1}$ für $n, k \in \mathbb{N}$,

so ist $f: \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ injektiv, also $|\mathbb{N} \times \mathbb{N}| \leq |\mathbb{N}|$, und damit nach Cantor-Bernstein $|\mathbb{N} \times \mathbb{N}| = |\mathbb{N}|$ (denn $i: \mathbb{N} \rightarrow \mathbb{N} \times \mathbb{N}$ mit $i(n) = (n, 0)$ für $n \in \mathbb{N}$ ist injektiv, also $|\mathbb{N}| \leq |\mathbb{N} \times \mathbb{N}|$).

Ebenso ist die Abbildung $g: \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ injektiv mit

$g(a, b) = p_0^{a+1} p_1^{b+1} = 2^{a+1} 3^{b+1}$ für $a, b \in \mathbb{N}$.

Analog erhält man für beliebige $n \geq 2$ injektive Funktionen $g_n: \mathbb{N}^{n+1} \rightarrow \mathbb{N}$. Wir setzen

$g_n(a_0, \dots, a_n) = p_0^{a_0+1} \dots p_n^{a_n+1}$ für $a_0, \dots, a_n \in \mathbb{N}$.

Bijektionen $g: \mathbb{N}^n \rightarrow \mathbb{N}$ werden auch als *Kodierungen* bezeichnet. Die natürliche Zahl $g(a_1, \dots, a_n)$ ist dann der Code für das „Wort“ $(a_1, \dots, a_n)$, und die Umkehrfunktion liefert die Dekodierung. Es gibt auch konkrete Kodierungen, die auf Worten variabler Länge operieren (siehe hierzu die Übung unten). In der Logik spielen sie eine große Rolle, da man mit ihrer Hilfe innerhalb der Zahlentheorie über Zeichenketten, und damit in der Zahlentheorie über eine formalisierte Mathematik reden kann.

Abzählbare Vereinigungen

Mit Hilfe einer Paarungsfunktion auf $\mathbb{N}^2$ können wir jetzt leicht einen sehr starken Satz beweisen.

Satz *(eine abzählbare Vereinigung abzählbarer Mengen ist abzählbar)*

Seien A_n abzählbare Mengen für $n \in \mathbb{N}$, und sei $A = \bigcup_{n \in \mathbb{N}} A_n$.

Dann ist A abzählbar.

Beweis

Für jedes $n \in \mathbb{N}$ sei

$f_n =$ „ein injektives $f_n : A_n \rightarrow \mathbb{N}$ “.

Sei $\pi : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ bijektiv (etwa die Cantorsche Paarungsfunktion).

Wir definieren $g : A \rightarrow \mathbb{N}$ durch:

$g(x) = \pi(f_n(x), n)$, wobei $n =$ „das kleinste $n \in \mathbb{N}$ mit $x \in A_n$ “.

– Dann ist $g : A \rightarrow \mathbb{N}$ injektiv, also $|A| \leq |\mathbb{N}|$.

Die Definition der Funktionen f_n geschieht wieder durch Auswahl. Für alle $n \in \mathbb{N}$ ist die Menge $M_n = \{g \mid g : A_n \rightarrow \mathbb{N} \text{ injektiv}\} \neq \emptyset$, denn A_n ist abzählbar. Für jedes $n \in \mathbb{N}$ wählen wir im Beweis ein $f_n \in M_n$.

Übung

Sei A eine Menge. Eine *endliche Folge in A* ist ein $f : \bar{n} \rightarrow A$ mit $n \in \mathbb{N}$.

Sei $F(A)$ die Menge aller endlichen Folgen in A .

- (i) Ist A abzählbar, so ist auch $F(A)$ abzählbar.
(Insbesondere ist also die „Menge aller Bücher“ abzählbar, selbst bei einem abzählbar unendlichen Zeichenvorrat A .)
- (ii) Ist A abzählbar, so ist auch die Menge aller endlichen Teilmengen von A abzählbar.
- (iii) Geben Sie eine bijektive Funktion $g : F(\mathbb{N}) \rightarrow \mathbb{N}$ direkt an.
[Primzahlen!]

Der Autor erinnert sich gut, daß bei ihm selber die „Abzählbarkeit aller Bücher“ einen großen Eindruck hinterließ, als er ihr zum ersten Mal begegnete, und möchte deshalb bei dieser Idee noch ein wenig verweilen, und dem Leser hierzu noch etwas anbieten. Der Gedanke wird in vielen frühen Büchern ausgeführt, am schönsten aber vielleicht bei Hausdorff:

Hausdorff (1914): „... Auch auf außermathematische Dinge ist diese Betrachtung häufig übertragen worden. Aus einem ‚Alphabet‘, d.h. einer endlichen Menge von ‚Buchstaben‘, kann man eine abzählbare Menge endlicher Buchstabenkomplexe, d.h. ‚Worte‘ bilden, unter denen sich natürlich auch sinnlose wie abracadabra befinden. Nimmt man zu den Buchstaben weitere Elemente wie Interpunktionszeichen, Druckspatien, Ziffern, Noten usw. hinzu, so sieht man, daß auch die Menge aller Bücher, Kataloge,

Symphonien, Opern abzählbar ist und abzählbar bleiben würde, wenn man selbst abzählbar viele Zeichen (aber für jeden Komplex nur endlich viele) verwenden wollte. Beschränkt man dagegen, bei endlicher Zeichenzahl, die Komplexe auf eine Maximalzahl von Elementen, indem man etwa Worte von mehr als hundert Buchstaben und Bücher von mehr als einer Million Worten für unstatthaft erklärt, so werden diese Mengen endlich, und wenn man mit Giordano Bruno eine unendliche Menge von Weltkörpern annimmt, mit sprechenden, schreibenden und musizierenden Bewohnern, so folgt mit mathematischer Gewißheit, daß auf unendlich vielen dieser Weltkörper dieselbe Oper mit demselben Libretto, denselben Namen des Komponisten, des Textdichters, der Orchestermitglieder und Sänger aufgeführt werden muß.“

Ebenso verblüffend wäre es, wenn auf einem anderen Weltkörper ein gewisser Herr Puccini die Opern von Wagner geschrieben hätte, oder ein Herr Francis Bacon die Werke von William Shakespeare...

Es ist schwer vorstellbar, daß Felix Hausdorff bei dieser Passage nicht geschmunzelt haben sollte, etwa bei den Worten „für unsatthaft erklärt“. Andererseits steht hinter ihr ein „ernster“ Gedanke, der Hausdorff und seine Zeit sehr beschäftigt hatte, nämlich Nietzsches „ewige Wiederkunft des Gleichen“. Einige Anmerkungen hierzu findet der Leser in der Biographie von Hausdorff am Ende des zweiten Abschnitts.

Die rationalen Zahlen sind abzählbar

Satz (*Abzählbarkeit der rationalen Zahlen*)

① ist abzählbar.

Beweis

Für $q \in \mathbb{Q}$ seien $N(q) \in \mathbb{N}$ der Nenner und $Z(q) \in \mathbb{N}$ der Zähler des gekürzten Bruches $|q|$, wobei $|q| = q$ für $q \geq 0$, $|q| = -q$ für $q < 0$.

Wir setzen für $n \in \mathbb{N}$

$$A_n = \{q \in \mathbb{Q} \mid N(q) + Z(q) = n\}.$$

Dann ist jedes A_n abzählbar (sogar endlich), und

$$\mathbb{Q} = \bigcup_{n \in \mathbb{N}} A_n.$$

– Also ist $\mathbb{Q}$ abzählbar.

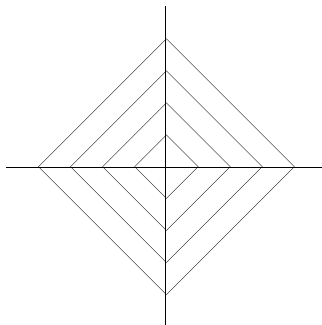
Hinter diesem Beweis steht die folgende „Spiralaufzählung“ des $\mathbb{Z}^2$ -Gitters.

Für $n \in \mathbb{N}$ setzen wir:

$$A'_n = \{ (a, b) \in \mathbb{R}^2 \mid |a| + |b| = n \}$$

Die Mengen A'_n sind Quadrate im $\mathbb{R}^2$ mit den Ecken $(n, 0)$, $(0, n)$, $(-n, 0)$, $(0, -n)$.

Vom Nullpunkt ausgehend können wir nun die Schnittpunkte dieser Quadrate mit dem $\mathbb{Z}^2$ -Gitter „spiralförmig“ aufzählen. Der Leser kann leicht die ersten Brüche einer solchen Aufzählung ermitteln.



Viele andere Arten von Aufzählungen von $\mathbb{Q}$ sind denkbar, etwa Abzählungen nach dem Maximum von Zähler und Nenner eines gekürzten Bruches, was dem um 45 Grad gedrehten Bild oben entsprechen würde. Beschränken wir uns auf rationale Zahlen q mit $0 \leq q \leq 1$, und ordnen wir bei gleichem Maximum nach aufsteigenden Zählern, so beginnt die Liste wie folgt:

$$\frac{0}{1} \quad \frac{1}{1} \quad \frac{1}{2} \quad \frac{1}{3} \quad \frac{2}{3} \quad \frac{1}{4} \quad \frac{3}{4} \quad \frac{1}{5} \quad \frac{2}{5} \quad \frac{3}{5} \quad \frac{4}{5} \quad \frac{1}{6} \quad \frac{5}{6} \quad \frac{1}{7} \quad \dots$$

In dieser algebraischen Form ist die Abzählbarkeit von $\mathbb{Q}$ einleuchtend. Überraschender ist sie, wenn man sich $\mathbb{Q}$ als Teilmenge der Zahlengeraden $\mathbb{R}$ vorstellt. Die Punkte von $\mathbb{Q}$ sind dicht gesät in $\mathbb{R}$:

Übung

$\mathbb{Q}$ ist dicht in $\mathbb{R}$, d. h. für alle $x, y \in \mathbb{R}$ mit $x < y$ existiert ein $q \in \mathbb{Q}$ mit $x < q < y$.

[Verwenden Sie z. B. Dezimalbruchentwicklungen von x, y , und interpolieren Sie eine Zahl q mit abbrechender Dezimalbruchentwicklung, also ein $q \in \mathbb{Q}$.]

Hausdorff (1914): „Die Äquivalenz der Menge der ganzen Zahlen mit der doch viel umfassenderen der rationalen Zahlen gehört mit zu den Tatsachen der Mengenlehre, die bei erster Bekanntschaft den Eindruck des Erstaunlichen, ja Paradoxen hervorrufen: namentlich wenn man das geometrische Bild (die Zuordnung zwischen Zahlen und Punkten der geraden Linie) vor Augen hat und sich einerseits die in endlichen Abständen isoliert liegenden ‚ganzzahligen‘ Punkte, andererseits die über die ganze Linie wie ein Staub von mehr als mikroskopischer Feinheit verteilten ‚rationalen‘ Punkte vergegenwärtigt.“

Aus der Dichtheit von $\mathbb{Q}$ in $\mathbb{R}$ zu schließen, daß auch $\mathbb{R}$ abzählbar ist, ist zwar auf den ersten Blick verführerisch, aber unstatthaft: Zwei geordnete Mengen als gleichgroß zu bezeichnen, wenn zwischen zwei verschiedenen Punkten der einen immer Punkte der anderen liegen, wäre kein guter Größenbegriff, und dieser Begriff hat mit der Existenz von Bijektionen i. a. nichts zu tun. (Der Autor erhält regelmäßig Zuschriften, in denen die Abzählbarkeit von $\mathbb{R}$ durch „ $\mathbb{Q}$ ist dicht in $\mathbb{R}$ “ bewiesen wird. Daher dieser vorbeugende Absatz.)

Die algebraischen Zahlen sind abzählbar

Definition (algebraische Zahlen)

Ein $x \in \mathbb{R}$ heißt *algebraisch*, falls x Nullstelle eines nichttrivialen Polynoms P mit ganzzahligen Koeffizienten ist, d. h. es existieren ein $n \in \mathbb{N}$ und $a_0, \dots, a_n \in \mathbb{Z}$, $a_n \neq 0$, mit:

$$P(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0 = 0.$$

Wir setzen $\mathbb{A} = \{ x \in \mathbb{R} \mid x \text{ ist algebraisch} \}$.

Übung

Es gilt $\mathbb{Q} \subseteq \mathbb{A}$. Weiter erhält man die gleiche Menge $\mathbb{A}$, wenn man Koeffizienten aus $\mathbb{Q}$ in den Polynomen zuläßt.

Es gilt $\mathbb{A} - \mathbb{Q} \neq \emptyset$. Denn sei x die positive Quadratwurzel aus 2, d. h. die Länge der Diagonalen eines Quadrats mit Seitenlänge 1. Dann gilt $x^2 - 2 = 0$, also ist x algebraisch. *Annahme*, $x = p/q \in \mathbb{Q}$. Dann gilt $(p/q)^2 = 2$, also $p^2 = 2q^2$. Der Faktor 2 hat in der Primfaktorzerlegung von p^2 einen geraden Exponenten – nämlich das Doppelte des entsprechenden Exponenten in der Primfaktorzerlegung von p –, dagegen hat er in der Zerlegung von $2q^2$ einen ungeraden Exponenten – nämlich das Doppelte des Exponenten der 2 in der Zerlegung von q plus 1 –, im *Widerspruch* zu $p^2 = 2q^2$ und der Eindeutigkeit der Primfaktorzerlegung.

Die Entdeckung der Existenz irrationaler Zahlen durch die Pythagoräer war für die – in geistigen Dingen – harmoniesüchtigen alten Griechen ein mathematischer Schock und bildet ein frühes Kapitel im Buch der allergischen Irritationen, die der erste Pollenflug neuer Zahlen anscheinend immer auslöst. Zu Zeiten Platons (427 – 347 v. Chr.) galt es dann schon als nicht besonders rühmlich, von der „Unverhältnismäßigkeit“ der Diagonalen eines Quadrat zu dessen Seite nichts zu wissen. In einer Zeit, in der allgemein angenommen wird, daß die Sterne ihr Licht von der Sonne haben, wäre es wohl unangebracht, über das Schattendasein der transfiniten Zahlen kulturpessimistisch zu lamentieren. Kommen wir also lieber zu unserem nächsten Satz, entdeckt und bewiesen von Cantor und Dedekind im Jahre 1873.

Satz (*Abzählbarkeit der algebraischen Zahlen*)

$\mathbb{A}$ ist abzählbar.

Beweis

Für $n \in \mathbb{N}$ sei

$$A_n = \{ x \in \mathbb{A} \mid x \text{ ist Nullstelle eines nichttrivialen Polynoms vom Grad } \leq n, \text{ dessen Koeffizienten } a \text{ alle } |a| \leq n \text{ erfüllen} \}$$

Da ein Polynom vom Grad n bekanntlich höchstens n reelle Nullstellen besitzt, ist jedes A_n abzählbar (sogar endlich). Es gilt

$$\mathbb{A} = \bigcup_{n \in \mathbb{N}} A_n.$$

– Also ist $\mathbb{A}$ abzählbar.

Georg Cantor über algebraische Zahlen und die Überabzählbarkeit des Kontinuums

„Unter einer reellen algebraischen Zahl wird allgemein eine reelle Zahlgröße ω verstanden, welche einer nicht identischen Gleichung von der Form genügt:

$$(1.) \quad a_0 \omega^n + a_1 \omega^{n-1} + \dots + a_n = 0,$$

wo $n, a_0, a_1, \dots, a_n$ ganze Zahlen sind; wir können uns hierbei die Zahlen n und a_0 positiv, die Koeffizienten $a_0, a_1, \dots, a_n$ ohne gemeinschaftlichen Teiler und die Gleichung (1.) irreduzibel denken; mit diesen Festsetzungen wird erreicht, daß nach den bekannten Grundsätzen der Arithmetik und Algebra die Gleichung (1.), welcher eine reelle algebraische Zahl genügt, eine völlig bestimmte ist; umgekehrt gehören bekanntlich zu jeder Gleichung von der Form (1.) höchstens so viele reelle algebraische Zahlen ω , welche ihr genügen, als ihr Grad n angibt. Die reellen algebraischen Zahlen bilden in ihrer Gesamtheit einen Inbegriff von Zahlgrößen, welcher mit (ω) bezeichnet werde; es hat derselbe, wie aus einfachen Betrachtungen hervorgeht, eine solche Beschaffenheit, daß in jeder Nähe irgendeiner gedachten Zahl α unendlich viele Zahlen aus (ω) liegen; um so auffallender dürfte daher für den ersten Anblick die Bemerkung sein, daß man den Inbegriff (ω) dem Inbegriffe aller ganzen positiven Zahlen v , welcher durch das Zeichen (v) angedeutet werde, eindeutig zuordnen kann, so daß zu jeder algebraischen Zahl ω eine bestimmte ganze positive Zahl v und umgekehrt zu jeder ganzen positiven Zahl v eine völlig bestimmte reelle algebraische Zahl ω gehört, daß also, um mit anderen Worten dasselbe zu bezeichnen, der Inbegriff (ω) in der Form einer unendlichen gesetzmäßigen Reihe:

$$(2.) \quad \omega_1, \omega_2, \dots, \omega_v, \dots$$

gedacht werden kann, in welcher sämtliche Individuen von (ω) vorkommen und ein jedes von ihnen sich an einer bestimmten Stelle in (2.), welche durch den zugehörigen Index gegeben ist, befindet. Sobald man ein Gesetz gefunden hat, nach welchem eine solche Zuordnung gedacht werden kann, läßt sich dasselbe nach Willkür modifizieren; es wird daher genügen, wenn ich in § 1 denjenigen Anordnungsmodus mitteile, welcher, wie mir scheint, die wenigsten Umstände in Anspruch nimmt.

Um von dieser Eigenschaft des Inbegriffs aller reellen algebraischen Zahlen eine Anwendung zu geben, füge ich zu § 1 den § 2 hinzu, in welchem ich zeige, daß, wenn eine beliebige Reihe reeller Zahlgrößen von der Form (2.) vorliegt, man in jedem vorgegebenen Intervalle $(\alpha \dots \beta)$ Zahlen η bestimmen kann, welche nicht in (2.) enthalten sind; kombiniert man die Inhalte dieser beider Paragraphen, so ist damit ein neuer Beweis des zuerst von *Liouville* bewiesenen Satzes gegeben, daß es in jedem Intervalle $(\alpha \dots \beta)$ unendlich viele transzendente, d. h. nicht algebraische reelle Zahlen gibt. Ferner stellt sich der Satz in § 2 als der Grund dar, warum Inbegriffe reeller Zahlgrößen, die ein sogenanntes Kontinuum bilden (etwa die sämtlichen reellen Zahlen, welche ≥ 0 und ≤ 1 sind), sich nicht eindeutig auf den Inbegriff (v) beziehen lassen; so fand ich den deutlichen Unterschied zwischen einem sogenannten Kontinuum und einem Inbegriffe von der Art der Gesamtheit aller reellen algebraischen Zahlen.“

(Georg Cantor 1874, „Über eine Eigenschaft des Inbegriffes aller reellen algebraischen Zahlen“)

8. Überabzählbare Mengen

Definition (*überabzählbare Mengen*)

Eine Menge M heißt *überabzählbar*, falls M nicht abzählbar ist.

Offenbar sind äquivalent:

- (i) M ist überabzählbar.
- (ii) $\text{non}(|M| \leq |\mathbb{N}|)$, d.h. es existiert kein injektives $f: M \rightarrow \mathbb{N}$.
- (iii) Es existiert kein surjektives $f: \mathbb{N} \rightarrow M$.
- (iv) $|\mathbb{N}| < |M|$.

Cantors Diagonalargument

Bislang haben wir nur abzählbare Mengen kennengelernt. Jetzt aber zeigt sich, daß die neben den natürlichen Zahlen wichtigste Struktur der Mathematik, die reellen Zahlen $\mathbb{R}$, nicht abzählbar ist. Cantor hat hierfür zwei Beweise gefunden, der erste benutzt die Vollständigkeit von $\mathbb{R}$, der zweite die Dezimaldarstellung einer reellen Zahl. Wir besprechen zunächst den berühmten zweiten Beweis. Das dabei auftauchende Argument einer Diagonalisierung wird in der Mengenlehre und in der Logik heute an verschiedenen Stellen benutzt.

Cantor formuliert die Frage zum ersten Mal in einem Brief an Dedekind gegen Ende des Jahres 1873. Der Brief zeigt, daß Cantor die Abzählbarkeit der rationalen Zahlen und auch die der Menge der endlichen Folgen natürlicher Zahlen zu diesem Zeitpunkt bereits kannte.

Cantor an Dedekind am 29.11.1873:

„Hochgeehrter Herr Kollege!

Gestatten Sie mir, Ihnen eine Frage vorzulegen, die für mich ein gewisses theoretisches Interesse hat, die ich mir aber nicht beantworten kann; vielleicht können Sie es, und sind so gut, mir darüber zu schreiben, es handelt sich um folgendes.

Man nehme den Inbegriff aller positiven ganzzahligen Individuen n und bezeichne ihn mit (n) ; ferner denke man sich etwa den Inbegriff aller positiven reellen Zahlgrößen x und bezeichne ihn mit (x) ; so ist die Frage einfach die, ob sich (n) dem (x) so zuordnen lasse, daß zu jedem Individuum des einen Inbegriffes ein und nur eines des andern gehört? Auf den ersten Anblick sagt man sich, nein es ist nicht möglich, denn (n) besteht aus diskreten Teilen, (x) aber bildet ein Kontinuum; nur ist mit diesem Einwande aber nichts gewonnen und so sehr ich mich auch zu der Ansicht neige, daß (n) und (x) keine eindeutige Zuordnung gestatten, so kann ich doch den Grund nicht finden und um den ist es mir zu tun, vielleicht ist es ein sehr einfacher. –

Wäre man nicht auch auf den ersten Anblick geneigt zu behaupten, daß sich (n) nicht eindeutig zuordnen lasse dem Inbegriffe (p/q) aller positiven rationalen Zahlen p/q ? Und dennoch ist es nicht schwer zu zeigen, daß sich (n) nicht nur diesem Inbegriffe, sondern noch dem allgemeineren

$$(a_{n_1, n_2, \dots, n_v})$$

eindeutig zuordnen läßt, wo $n_1, n_2, \dots, n_v$ unbeschränkte positive ganzzahlige Indizes in beliebiger Zahl v sind.

Mit bestem Gruße
Ihr ergebenster
G. Cantor

Den „Grund“ konnte Cantor Dedekind bereits etwa eine Woche nach der Formulierung des Problems mitteilen: Das Datum des entsprechenden Briefes an Dedekind, der 7. 12. 1873, wird häufig als der Geburtstag der Mengenlehre bezeichnet. Einen weiteren Beweis der Überabzählbarkeit von $\mathbb{R}$ fand Cantor später. Er trug ihn 1891 auf der ersten Jahresversammlung der von ihm mitbegründeten Deutschen Mathematiker-Vereinigung vor. Wir bringen hier zuerst diesen späteren Beweis, der zu einem Klassiker der Mathematik geworden ist.

Satz (*Satz von Cantor über die Überabzählbarkeit der reellen Zahlen*)

Die Menge $\mathbb{R}$ der reellen Zahlen ist überabzählbar.

Beweis

Es genügt zu zeigen: Es existiert kein $f: \mathbb{N} \rightarrow \mathbb{R}$ surjektiv.

Sei hierzu $f: \mathbb{N} \rightarrow \mathbb{R}$ beliebig.

Wir finden ein $x \in \mathbb{R}$ mit $x \notin \text{rng}(f)$ wie folgt.

Für $n \in \mathbb{N}$ schreiben wir $f(n)$ in kanonischer unendlicher

Dezimaldarstellung [mit $0 = 0,000 \dots$].

Sei also:

$$f(0) = z_0, a_{0,0} \ a_{0,1} \ a_{0,2} \dots$$

$$f(1) = z_1, a_{1,0} \ a_{1,1} \ a_{1,2} \dots$$

$$f(2) = z_2, a_{2,0} \ a_{2,1} \ a_{2,2} \dots$$

$$f(3) = z_3, a_{3,0} \ a_{3,1} \ a_{3,2} \dots$$

...

$$f(n) = z_n, a_{n,0} \ a_{n,1} \ a_{n,2} \dots$$

...

Wir definieren $x = 0, b_0 \ b_1 \ b_2 \dots \in \mathbb{R}$ durch

$$b_n = \begin{cases} 1, & \text{falls } a_{n,n} = 2, \\ 2, & \text{falls } a_{n,n} \neq 2. \end{cases}$$

Dann ist $x = 0, b_0 \ b_1 \ b_2 \dots$ in kanonischer Dezimaldarstellung.

Für jedes $n \in \mathbb{N}$ gilt nun aber $x \neq f(n)$, denn die n -ten Nachkommastellen der kanonischen Dezimaldarstellungen von x und $f(n)$ sind verschieden, und die kanonische Dezimaldarstellung einer reellen Zahl ist eindeutig.

– Also $x \notin \text{rng}(f)$, und damit ist f nicht surjektiv.

Übung

Warum ist es ungünstig, 0 und 1 in der Definition von b_n zu verwenden an Stelle von 1 und 2?

[Wir setzen $f(0) = 0,099999 \dots$,
 $f(1) = 0,011111 \dots$,
 $f(2) = 0,001111 \dots$,
 $f(3) = 0,000111 \dots$, allgemein
 $f(n) = 1/9 \cdot 1/10^n$ für $n \geq 1$.

Dann ist die mittels 0 und 1 definierte Diagonalisierung nicht in kanonischer Darstellung und gleich $f(0)$, also im Wertebereich von f .]

Ein Korollar zu diesem Satz betrifft die Existenz von transzendenten Zahlen. Dies sind reelle Zahlen, die sich nicht als Nullstellen von Polynomen mit rationalen Koeffizienten darstellen lassen:

Definition (*transzendente Zahlen*)

Sei x eine reelle Zahl.

x heißt *transzendent*, wenn x nicht algebraisch ist.

$\mathbb{T}$ sei die Menge der transzendenten Zahlen.

Der Nachweis der Transzendenz einer bestimmten Zahl ist im allgemeinen ein schwieriges Problem. 1851 konnte Joseph Liouville (1809 – 1882) zeigen, daß die Zahl

$0,1100010000000000000000010 \dots$

transzendent ist, wobei die m -te Nachkommastelle dieser Zahl genau dann gleich 1 ist, wenn $m = n!$ für ein $n \geq 1$. Liouville wollte die Transzendenz der Eulerschen Zahl $e = \sum_{n \geq 0} 1/n!$ beweisen, was dann erst Charles Hermite (1822 – 1902) 1872 gelang. Auf den Arbeiten von Hermite aufbauend bewies Lindemann 1882, daß die Kreiszahl π transzendent ist. (Cantor hat diese Arbeit referiert.) Der nächste große Schritt war die Lösung des siebenten Hilbertschen Problems durch Alexander Gelfond (1906 – 1968) und Theodor Schneider (1911 – 1988), die im Jahre 1934 unabhängig voneinander zeigten [siehe etwa Siegel 1949]:

Ist a algebraisch, $a > 0$, $a \neq 1$, und ist b irrational und algebraisch, so ist a^b transzendent.

Ist der Nachweis im Einzelfall schwierig, so zeigt der Satz von Cantor doch, daß „fast alle“ reellen Zahlen transzendent sind:

Korollar (*die transzendenten Zahlen sind überabzählbar*)

$\mathbb{T}$ ist überabzählbar. Genauer gilt: $\mathbb{R} - \mathbb{T}$ ist abzählbar.

Beweis

Es gilt $\mathbb{R} = \mathbb{A} \cup \mathbb{T}$. Da die Menge $\mathbb{A}$ der algebraischen Zahlen abzählbar ist, ist notwendig $\mathbb{T}$ überabzählbar, denn die Vereinigung

– zweier abzählbarer Mengen ist abzählbar.

Übung

Seien M_n überabzählbare Mengen für $n \in \mathbb{N}$.

Für alle $n \in \mathbb{N}$ sei $M_n - M_{n+1}$ abzählbar.

Dann ist $\bigcap_{n \in \mathbb{N}} M_n$ überabzählbar.

Cantors ursprünglicher Beweis der Überabzählbarkeit von $\mathbb{R}$

Für sich von Interesse ist auch der erste Cantorsche Beweis der Überabzählbarkeit von $\mathbb{R}$, der die Vollständigkeit der reellen Zahlen benutzt: Jede nichtleere nach oben beschränkte Teilmenge von $\mathbb{R}$ hat ein Supremum. Genauer bedeutet dies das folgende:

Sei $X \subseteq \mathbb{R}$ und $a \in \mathbb{R}$. Wir schreiben $X \leq a$, falls $x \leq a$ gilt für alle $x \in X$. Die *Vollständigkeit* von $\mathbb{R}$ lautet nun:

Sei $X \subseteq \mathbb{R}$ nichtleer und es existiere ein $a \in \mathbb{R}$ mit $X \leq a$.

Dann existiert ein eindeutiges $b \in \mathbb{R}$ mit:

(i) $X \leq b$, (ii) für alle $c \in \mathbb{R}$ mit $X \leq c$ gilt $b \leq c$.

b heißt das *Supremum* oder *die kleinste obere Schranke* von X , in Zeichen $b = \sup(X)$.

Anschaulich: Man wählt zu einer nach oben beschränkten Teilmenge X von $\mathbb{R}$ ein a mit $X \leq a$. Nun wird diese obere Schranke a von X solange nach unten verschoben, bis sie X von oben berührt. Der Berührungspunkt ist gerade $\sup(X)$. (Sowohl $\sup(X) \in X$ als auch $\sup(X) \notin X$ sind möglich.)

Die Vollständigkeit unterscheidet die reellen Zahlen wesentlich von den rationalen Zahlen und wird in der Analysis an allen Ecken und Enden gebraucht. (So gilt etwa der Zwischenwertsatz nicht für stetige Funktionen $f: \mathbb{Q} \rightarrow \mathbb{Q}$: Sei $f(x) = x^2 - 2$ für $x \in \mathbb{Q}$. Dann ist $f: \mathbb{Q} \rightarrow \mathbb{Q}$ stetig, $f(0) = -2 < 0$, $f(2) = 2 > 0$, aber es gibt kein $x \in \mathbb{Q}$ mit $f(x) = 0$, denn die Quadratwurzel aus 2 ist irrational.)

Den folgenden Beweis der Überabzählbarkeit von $\mathbb{R}$ hat Cantor 1874 veröffentlicht. Er ist eine vereinfachte Version des Beweises, den Cantor Dedekind am 7.12.1873 brieflich mitgeteilt hatte.

Der ursprüngliche Beweis der Überabzählbarkeit der reellen Zahlen

Sei $f: \mathbb{N} \rightarrow \mathbb{R}$ beliebig. Wir suchen ein $x^* \in \mathbb{R}$ mit $x^* \notin \text{rng}(f)$.

Wir setzen $x_n = f(n)$ für $n \in \mathbb{N}$. Es gilt also $\text{rng}(f) = \{x_n \mid n \in \mathbb{N}\}$.

Wir definieren nun rekursiv solange möglich:

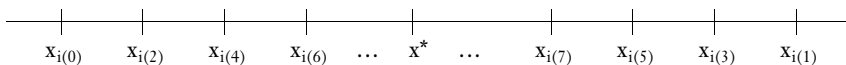
$i(0) = 0$,

$i(1) = \text{„das kleinste } k \in \mathbb{N} \text{ mit } x_0 < x_k\text{“}$.

$i(n) = \text{„das kleinste } k \in \mathbb{N} \text{ mit:}$

$x_k \text{ liegt zwischen } x_{i(n-2)} \text{ und } x_{i(n-1)}\text{“}$ für alle $n \geq 2$,

wobei wir sagen: $x \in \mathbb{R}$ *liegt zwischen* $a, b \in \mathbb{R}$, falls $a < x < b$ oder $b < x < a$.



Ist $i(n)$ nicht definiert für ein $n \in \mathbb{N}$, so ist offenbar $\text{rng}(f) \neq \mathbb{R}$.

[Es existieren dann sogar $a, b \in \mathbb{R}$, $a \neq b$, mit der Eigenschaft:
für alle x zwischen a und b gilt $x \notin \text{rng}(f)$.]

Die Funktion f läßt dann sogar ein ganzes Intervall aus.]

Sei also $i(n)$ definiert für alle $n \in \mathbb{N}$.

Nach Konstruktion ist $i(n) < i(n+1)$ für alle $n \in \mathbb{N}$ (!), und es gilt

$$x_{i(0)} < x_{i(2)} < x_{i(4)} < \dots < x_{i(5)} < x_{i(3)} < x_{i(1)}.$$

Wir setzen nun:

$$x^* = \sup \{ x_{i(n)} \mid n \text{ gerade} \} \in \mathbb{R}.$$

Dann liegt x^* zwischen $x_{i(n)}$ und $x_{i(n+1)}$ für alle $n \in \mathbb{N}$.

Weiter ist $x^* \notin \text{rng}(f)$:

Annahme, es gibt ein $k \in \mathbb{N}$ mit $x^* = x_k$. Sei dann n derart, daß

$$i(n-2) < k < i(n-1).$$

— Dann ist $i(n) \leq k$ nach Definition von $i(n)$, also $i(n) < i(n-1)$, *Widerspruch*.

Wir können die Konstruktion dieses Beweises für eine Bijektion $f: \mathbb{N} \rightarrow \mathbb{Q}$ zwischen den natürlichen und den rationalen Zahlen durchführen; dann sind alle $i(n)$ für $n \in \mathbb{N}$ definiert, denn für zwei verschiedene rationale Zahlen a und b existiert immer eine rationale Zahl x zwischen a und b . $x^* = \sup \{ x_{i(n)} \mid n \text{ gerade} \}$ ist dann notwendig nicht im Wertebereich $\mathbb{Q}$ von f , also notwendig eine irrationale Zahl. Starten wir mit einer Bijektion $f: \mathbb{N} \rightarrow \mathbb{A}$ zwischen den natürlichen und den algebraischen Zahlen, erhalten wir eine transzendente Zahl x^* .

Einfache Folgerungen

Der ursprüngliche Beweis von Cantor zeigt sofort:

Korollar (*jedes reelle Intervall ist überabzählbar*)

Seien $x, y \in \mathbb{R}$ mit $x < y$.

Dann ist $]x, y[= \{ z \in \mathbb{R} \mid x < z < y \}$ überabzählbar.

Übung

Zeigen Sie die Überabzählbarkeit jedes reellen Intervalls durch eine Modifikation des Diagonalarguments.

Wir geben noch einen weiteren Beweis mit Hilfe von Translationen.

Weiterer Beweis des Korollars

Für ein $A \subseteq \mathbb{R}$ und ein $z \in \mathbb{R}$ definieren wir $A + z = \{x + z \mid x \in A\}$.

$A + z$ heißt *die Translation von A um z*. Anschaulich wird die Menge A um den Wert z verschoben. Offenbar gilt $]x, y[+ z =]x + z, y + z[$.

Weiter gilt $|A| = |A + z|$, denn $f: A \rightarrow A + z$ mit $f(x) = x + z$ ist bijektiv.

Sei nun $]x, y[$ ein beliebiges Intervall mit $x < y$, $x, y \in \mathbb{R}$.

Annahme, $]x, y[$ ist abzählbar.

Es gilt aber $\mathbb{R} = \bigcup \{]x, y[+ q \mid q \in \mathbb{Q} \} (!)$.

Da mit $]x, y[$ auch alle $]x, y[+ q$ abzählbar sind,
ist $\mathbb{R}$ abzählbare Vereinigung von abzählbaren Mengen.

– Also ist $\mathbb{R}$ abzählbar, *Widerspruch!*

Stärker gilt:

Satz (Mächtigkeit reeller Intervalle)

Seien $a, b \in \mathbb{R}$, $a < b$. Dann gilt

$$|\mathbb{R}| = |]a, b[|.$$

Beweis

Sei $I =]-1, 1[$. Es genügt zu zeigen: $|\mathbb{R}| = |I|$, denn offene Intervalle verschiedener Länge kann man durch einfache Streckung bijektiv aufeinander abbilden:

Sind $a, b, c, d \in \mathbb{R}$ und $a < b$, $c < d$, so ist die Funktion

$f:]a, b[\rightarrow]c, d[$ bijektiv, wobei

$$f(y) = c + (y - a)/(b - a) \cdot (d - c) \text{ für } y \in]a, b[.$$

Wir definieren nun $g: I \rightarrow \mathbb{R}$ durch

$$g(x) = \begin{cases} 1/x - 1, & \text{falls } 0 < x < 1, \\ 0, & \text{falls } x = 0, \\ 1/x + 1, & \text{falls } -1 < x < 0. \end{cases}$$

– Dann ist $g: I \rightarrow \mathbb{R}$ bijektiv.

Es gibt viele konkrete Funktionen, die Intervalle bijektiv auf ganz $\mathbb{R}$ abbilden, z. B. ist die Tangensfunktion $\tan:]-\pi/2, \pi/2[\rightarrow \mathbb{R}$ bijektiv (und stetig).

Mit dem Satz von Cantor-Bernstein folgt leicht, daß auch halboffene und abgeschlossene reelle Intervalle, die mehr als einen Punkt enthalten, gleichmächtig zu $\mathbb{R}$ sind; denn solche Intervalle enthalten ein offenes Intervall $]x, y[$ mit $x < y$. Einen direkteren Beweis gibt die folgende Übung.

Übung

Sei $I =]0, 1[$, $J =]0, 1] = \{x \in \mathbb{R} \mid 0 < x \leq 1\}$. Dann gilt $|I| = |J|$.

[Sei $H = \{1/n \mid n \in \mathbb{N}, n \geq 1\}$.

Betrachte $g = \{(1/n, 1/(n+1)) \mid n \geq 1\} \cup \text{id}_{J-H}$.

H ist zudem der Wertebereich des Orbits von 1 unter g .]

Gleichmächtig zu $\mathbb{R}$ ist weiter auch die Vereinigung von beliebig vielen Intervallen, falls mindestens ein Intervall der Vereinigung nichttrivial ist.

Übung

Es gilt $|\mathbb{R}| = |\mathbb{R} \times \mathbb{N}|$.

[Sei $I = [0, 1[= \{x \in \mathbb{R} \mid 0 \leq x < 1\}$. Es gilt $|I| = |\mathbb{R}|$, und

$f: I \times \mathbb{N} \rightarrow \mathbb{R}$ mit $f(x, n) = n + x$ ist injektiv.]

Für $r \in \mathbb{R}$, $0 < r$, sei K_r die Kreislinie in $\mathbb{R}^2 = \mathbb{R} \times \mathbb{R}$ mit Radius r um den Nullpunkt. Für beliebige $r, r' \in \mathbb{R}$ mit $0 < r < r'$ erhält man eine Bijektion zwischen K_r und $K_{r'}$, indem man die Punkte der Kreislinien aufeinander abbildet, die auf den gleichen im Nullpunkt beginnenden Halbstrahlen liegen (d. h. die Punkte der Kreislinien mit gleichem Winkel zur x -Achse werden einander zugeordnet). Nach obiger Übung ist also $|\mathbb{R}| = |\bigcup_{r>0, r \in \mathbb{A}} K_r|$, da $|\mathbb{R}| = |K_1|$. Wir werden im nächsten Kapitel zeigen, daß sogar der gesamte $\mathbb{R}^2$ gleichmächtig zu $\mathbb{R}$ ist!

Wir schließen hier noch eine Übung an, die zeigt, wie dünn $\mathbb{Q} \times \mathbb{Q}$ in $\mathbb{R} \times \mathbb{R}$ gesät ist. Die schöne Idee, Kreisbögen zu verwenden, stammt von Cantor.

Übung

Seien $x, y \in \mathbb{R}^2$, $x \neq y$. Sei $\mathbb{P} = \mathbb{R}^2 - \mathbb{Q}^2$.

Ein *Kreisbogen zwischen x und y* ist ein x und y verbindendes Segment eines Kreises im $\mathbb{R}^2$; per Konvention rechnen wir hier x und y nicht zu einem Kreisbogen zwischen x und y mit dazu.

Zeigen Sie: Es gibt einen Kreisbogen B zwischen x und y mit $B \subseteq \mathbb{P}$.

[Es gibt überabzählbar viele paarweise disjunkte Kreisbögen zwischen x und y .

Die Aussage gilt allgemein für $\mathbb{P} = \mathbb{R}^2 - A$ mit $A \subseteq \mathbb{R}^2$ abzählbar.]

Wir können also verschiedene Punkte x und y im $\mathbb{R}^2$ durch eine stetige Linie – sogar einen Kreisbogen – derart verbinden, daß die Verbindungslinie an allen Punkten der Menge $\mathbb{Q} \times \mathbb{Q}$ vorbeiläuft – obwohl die Menge $\mathbb{Q} \times \mathbb{Q}$ dicht in $\mathbb{R}^2$ liegt (d. h. für alle $x \in \mathbb{R}^2$ und alle $r > 0$ ist $\mathbb{Q}^2 \cap K(x, r) \neq \emptyset$, wobei $K(x, r) \subseteq \mathbb{R}^2$ hier den Vollkreis um x mit Radius r bezeichnet).

Cantor (1882b): „Was die abzählbaren *Punktmen*gen anbetrifft, so bieten sie eine merkwürdige Erscheinung dar, welche ich im Folgenden zum Ausdruck bringen möchte. Betrachten wir irgendeine Punktmenge (M), welche innerhalb eines n -dimensionalen stetig zusammenhängenden Gebietes A *überalldicht* verbreitet ist und die Eigenschaft der Abzählbarkeit besitzt, so daß die zu (M) gehörigen Punkte sich in der Reihenform:

$$M_1, M_2, \dots, M_v, \dots$$

vorstellen lassen; als Beispiel diene die Menge aller derjenigen Punkte unseres dreidimensionalen Raumes, deren Koordinaten in bezug auf ein orthogonales Koordinatensystem x, y, z alle drei *algebraische* Zahlenwerte haben. Denkt man sich aus dem Gebiete A die abzählbare Punktmenge (M) entfernt und das alsdann übrig gebliebene Gebiet mit $\mathfrak{A}$ bezeichnet, so besteht der merkwürdige Satz, daß für $n \geq 2$ das Gebiet $\mathfrak{A}$ *nicht aufhört stetig zusammenhängend zu sein*, daß mit anderen Worten je zwei Punkte N und N' des Gebietes $\mathfrak{A}$ immer verbunden werden können durch eine *stetige Linie*, welche mit allen ihren Punkten dem Gebiete $\mathfrak{A}$ angehört, so daß auf ihr kein einziger Punkt der Menge (M) liegt.“

Die Klärung des allgemeinen mathematischen Raumbegriffs mit Konzepten wie „zusammenhängend“, „wegzusammenhängend“, usw., geschah erst Anfang des 20. Jahrhunderts durch die aus der Mengenlehre hervorgehende Topologie. Felix Hausdorff ist hier eine zentrale Figur, er definierte 1914 allgemeine *topologische* und speziellere *metrische Räume*. Cantor war an Fragen des mathematischen Raumbegriffs und der Inhaltsmessung von Punktmengen in einem Raum sehr interessiert, und hat nicht nur durch seine allgemeine mengentheoretische Sprache, sondern speziell durch seine Untersuchung von „Punktmannigfaltigkeiten“ die allgemeine Raum- und Maßtheorie initiiert. Wir werden im zweiten Abschnitt Cantors Analyse von Punktmengen im Raum der reellen Zahlen ausführlich behandeln.

Subtraktion einer abzählbaren Menge

Wir wissen bereits, daß durch Entfernen einer abzählbaren Teilmenge die Überabzählbarkeit einer Menge nicht verändert wird. Wir zeigen nun stärker, daß die Kardinalität einer Menge durch Entfernung abzählbar vieler Elemente nicht verändert wird:

Satz (*Subtraktion abzählbarer Mengen*)

Sei M eine überabzählbare Menge. Weiter sei $A \subseteq M$ abzählbar.
Dann gilt $|M - A| = |M|$.

Beweis

Wir nehmen zunächst an, daß A abzählbar unendlich ist.

Sei $f: \mathbb{N} \rightarrow A$ bijektiv, und sei $x_n = f(n)$ für $n \in \mathbb{N}$.

$M - A$ ist eine unendliche Menge (sogar überabzählbar).

Sei also $B \subseteq M - A$ eine abzählbar unendliche Teilmenge von $M - A$.

Sei $g: \mathbb{N} \rightarrow B$ bijektiv und $y_n = g(n)$ für alle $n \in \mathbb{N}$.

Wir definieren $h: M - A \rightarrow M$ durch:

$$h(z) = \begin{cases} y_n, & \text{falls } z = y_{2n}, \\ x_n, & \text{falls } z = y_{2n+1}, \\ z, & \text{falls } z \notin B. \end{cases}$$

Dann ist $h: M - A \rightarrow M$ bijektiv, also $|M| = |M - A|$.

Sei nun A endlich, $f: \bar{m} \rightarrow A$ bijektiv für ein $m \in \mathbb{N}$, $x_n = f(n)$ für $n < m$.

Seien B, g, y_n für $n \in \mathbb{N}$ wie eben definiert.

Wir definieren $h: M - A \rightarrow M$ durch:

$$h(z) = \begin{cases} y_{n-m}, & \text{falls } z = y_n, n \geq m \\ x_n, & \text{falls } z = y_n, 0 \leq n < m \\ z, & \text{falls } z \notin B. \end{cases}$$

– Dann ist $h: M - A \rightarrow M$ bijektiv, also $|M| = |M - A|$.

Korollar

Es gilt $|\mathbb{R}| = |\mathbb{R} - \mathbb{Q}| = |\mathbb{T}|$.

Die reellen Zahlen sind also gleichmächtig zu den irrationalen Zahlen und stärker sogar gleichmächtig zu den transzendenten Zahlen.

Das „Reißverschluß“-Argument des Beweises im Subtraktionssatz stammt von Cantor. In einer Arbeit von 1878 benötigt er die Gleichung $|I| = |I - \mathbb{Q}|$, wobei $I = [0, 1] \subseteq \mathbb{R}$ ist. Nach einem vierseitigen umständlichen Beweis dieses Resultates (der Satz von Cantor-Bernstein stand ihm damals noch nicht zur Verfügung) gibt er schließlich einen zweiten Beweis nach obiger gerade-ungerade-Aufspaltung einer abzählbar unendlichen Teilmenge B von $I - \mathbb{Q}$. Auf den komplizierten Beweis, der ihn viel Mühe gekostet hatte, wollte er nicht verzichten, „weil die Hilfssätze (F.), (G.), (H.), (J.), welche bei der komplizierten Beweisführung gebraucht werden, an sich von Interesse sind“. Das Interesse an diesen Hilfssätzen hält sich in Grenzen. Das rückblickende Finden einfacher Argumente ist für die Mitwelt meistens ein Glücksfall, für einen Forscher selber bedeutet es oft, daß viel vorangehende Arbeit überflüssig wird. Und es kann schwerfallen, die erste Wegbeschreibung zu einem neuen Resultat einfach in den Papierkorb zu werfen.

Den Schluß dieses Kapitels bildet Cantors erster Beweis der Überabzählbarkeit der reellen Zahlen, in seinen eigenen Worten vorgetragen.

Georg Cantor über die Überabzählbarkeit des Kontinuums

„§. 2.

Wenn eine nach irgend einem Gesetze gegebene unendliche Reihe von einander verschiedener reeller Zahlgrößen:

$$(4.) \quad \omega_1, \omega_2, \dots, \omega_v, \dots$$

vorliegt, so läßt sich in jedem vorgegebenen Intervalle $(\alpha \dots \beta)$ eine Zahl η (und folglich unendlich viele solcher Zahlen) bestimmen, welche in der Reihe (4.) nicht vorkommt; dies soll nun bewiesen werden.

Wir gehen zu dem Ende von dem Intervalle $(\alpha \dots \beta)$ aus, welches uns beliebig vorgegeben sei, und es sei $\alpha < \beta$; die ersten beiden Zahlen unserer Reihe (4.), welche im Inneren dieses Intervalls (mit Ausschluß der Grenzen) liegen, mögen mit α' , β' bezeichnet werden, und es sei $\alpha' < \beta'$; ebenso bezeichne man in unserer Reihe die ersten beiden Zahlen, welche im Inneren von $(\alpha' \dots \beta')$ liegen, mit α'' , β'' , und es sei $\alpha'' < \beta''$, und nach demselben Gesetze bilde man ein folgendes Intervall $(\alpha''' \dots \beta''')$ u. s. w. Hier sind also α' , $\alpha'' \dots$ der Definition nach bestimmte Zahlen unserer Reihe (4.), deren Indizes im fortwährend Steigen sich befinden, und das gleiche gilt von den Zahlen β' , $\beta'' \dots$; ferner nehmen die Zahlen α' , $\alpha'' \dots$ ihrer Größe nach fortwährend zu, die Zahlen β' , β'' nehmen ihrer Größe nach fortwährend ab; von den Intervallen $(\alpha \dots \beta)$, $(\alpha' \dots \beta')$, $(\alpha'' \dots \beta'')$, ... schließt ein jedes alle auf dasselbe folgenden ein. – Hierbei sind nun zwei Fälle denkbar.

Entweder die Anzahl der so gebildeten Intervalle ist endlich; das letzte von ihnen sei $(\alpha^{(v)} \dots \beta^{(v)})$; da im Inneren desselben höchstens eine Zahl der Reihe (4.) liegen kann, so kann eine Zahl η in diesem Intervalle angenommen werden, welche nicht in (4.) enthalten ist, und es ist somit der Satz für diesen Fall bewiesen. –

Oder die Anzahl der gebildeten Intervalle ist unendlich groß; dann haben die Größen α , α' , α'' , ..., weil sie fortwährend ihrer Größe nach zunehmen ohne ins Unendliche zu

wachsen, einen bestimmten Grenzwert α^∞ ; ein gleiches gilt für die Größen $\beta, \beta', \beta'', \dots$, weil sie fortwährend ihrer Größe nach abnehmen, ihr Grenzwert sei β^∞ ; ist $\alpha^\infty = \beta^\infty$ (ein Fall, der bei dem Inbegriffe (ω) aller reellen algebraischen Zahlen stets eintritt), so überzeugt man sich leicht, wenn man nur auf die Definition der Intervalle zurückblickt, daß die Zahl $\eta = \alpha^\infty = \beta^\infty$ nicht in unserer Reihe enthalten sein kann *); ist aber $\alpha^\infty < \beta^\infty$, so genügt jede Zahl η im Inneren des Intervalles ($\alpha^\infty \dots \beta^\infty$) oder auch an den Grenzen desselben der gestellten Forderung, nicht in der Reihe (4.) enthalten zu sein. –

...

*) Wäre die Zahl η in unserer Reihe enthalten, so hätte man $\eta = \omega_p$, wo p ein bestimmter Index ist; dies ist aber nicht möglich, denn ω_p liegt nicht in Innern des Intervalls ($\alpha^{(p)} \dots \beta^{(p)}$), während die Zahl η ihrer Definition nach im Innern dieses Intervalls liegt.“

(Georg Cantor 1874, „Über eine Eigenschaft des Inbegriffes aller reellen algebraischen Zahlen“)

9. Mengen der Mächtigkeit der reellen Zahlen

Mehrdimensionale Kontinua

Für die natürlichen Zahlen haben wir gezeigt, daß $|\mathbb{N} \times \mathbb{N}| = |\mathbb{N}|$ gilt. Wir zeigen jetzt das analoge und in diesem Fall kontraintuitive Resultat für $\mathbb{R}$.

Cantor (Brief an Dedekind vom 5. 1. 1874):

„Hochgeehrter Herr Professor!

... Was die Fragen anbetrifft, mit denen ich in der letzten Zeit mich beschäftigt habe, so fällt mir ein, daß in diesem Gedankengange auch die folgende sich darbietet:

Läßt sich eine Fläche (etwa ein Quadrat mit Einschluß der Begrenzung) eindeutig auf eine Linie (etwa eine gerade Strecke mit Einschluß der Endpunkte) eindeutig beziehen, so daß zu jedem Punkte der Fläche ein Punkt der Linie und umgekehrt zu jedem Punkte der Linie ein Punkt der Fläche gehört?

Mir will es im Augenblick noch scheinen, daß die Beantwortung dieser Fragen, – obgleich man auch hier zum Nein sich so gedrängt sieht, daß man den Beweis dazu fast für überflüssig halten möchte, – große Schwierigkeiten hat. – ... “

Mehr als drei Jahre hat es gedauert, bis Cantor die überraschende Antwort auf das Problem fand. Brieflich teilt Cantor Dedekind am 20. 6. 1877 einen leicht fehlerhaften Beweis von $|I^n| = |I|$ mit, wobei $I = \{x \in \mathbb{R} \mid 0 \leq x \leq 1\}$ das abgeschlossene reelle Einheitsintervall ist und $n \geq 1$ beliebig; sein Argument zeigt lediglich $|I^n| \leq |I|$, was aber, wie Cantor betont, den Kern der Sache betrifft. In „Ein Beitrag zur Mannigfaltigkeitslehre“ (1878) konstruiert Cantor dann eine Bijektion von I^n nach I unter Verwendung von Kettenbrüchen. Heute ist der Beweis mit Hilfe des Satzes von Cantor-Bernstein oder einem Trick von Julius König (s. u.) einfach zu führen.

Satz (Satz von Cantor über die Mächtigkeit von $\mathbb{R} \times \mathbb{R}$)

Es gilt $|\mathbb{R} \times \mathbb{R}| = |\mathbb{R}|$.

Beweis

Es gilt $|\mathbb{R}| \leq |\mathbb{R} \times \mathbb{R}|$. Betrachte hierzu $i : \mathbb{R} \rightarrow \mathbb{R} \times \mathbb{R}$ mit $i(x) = (x, 0)$ für $x \in \mathbb{R}$. Dann ist i injektiv.

Es bleibt zu zeigen, daß $|\mathbb{R} \times \mathbb{R}| \leq |\mathbb{R}|$.

Sei $g : \mathbb{Z} \times \mathbb{Z} \rightarrow \mathbb{Z}$ bijektiv mit $g(0, 0) = 0$.

Sei $(x, y) \in \mathbb{R} \times \mathbb{R}$. Wir schreiben x und y in kanonischer Dezimaldarstellung:

$$x = c, a_0 a_1 a_2 \dots,$$

$$y = d, b_0 b_1 b_2 \dots$$

mit $c, d \in \mathbb{Z}$. Wir definieren nun $f: \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$ durch „Mischung“ der Nachkommastellen im Reißverschlußverfahren:

$$f(x, y) = g(c, d), a_0 b_0 a_1 b_1 a_2 b_2 \dots$$

Dann ist $f(x, y)$ in kanonischer Darstellung, und damit ist offenbar f injektiv.

- $(f(0, 0) = 0,000 \dots)$ ist in kanonischer Darstellung wegen $g(0, 0) = 0$.)

Übung

Die Abbildung f im obigen Beweis ist nicht surjektiv.

Genauer gilt: $\mathbb{R} - \text{rng}(f)$ ist überabzählbar.

Da jedes Intervall $]a, b[$, $a, b \in \mathbb{R}$, $a < b$, die Mächtigkeit von $\mathbb{R}$ hat, folgt: Jedes reelle Intervall läßt sich bijektiv auf die ganze Ebene $\mathbb{R}^2$ abbilden.

Das Ergebnis $|\mathbb{R} \times \mathbb{R}| = |\mathbb{R}|$ hat zur Zeit seiner Entdeckung große Irritationen hervorgerufen, auch bei Cantor selbst, der in einem Brief an Dedekind das Französische zu Hilfe ruft: „je le vois, mais je ne crois pas“ [„ich sehe es, aber ich glaube es nicht“].

Die Gleichung $|\mathbb{R}^2| = |\mathbb{R}|$ erlaubt uns, Teilmengen der Ebene als Teilmengen der Geraden anzusehen – wir wählen ein bijektives $f: \mathbb{R}^2 \rightarrow \mathbb{R}$ und setzen $B = f^{-1}A$ für $A \subseteq \mathbb{R}^2$. Allerdings werden bei diesem Übergang von A zu B wesentliche Strukturen von B zerstört; die Abbildung f ist unstetig, sie schüttelt gewissermaßen den $\mathbb{R}^2$ völlig durcheinander, um ihn danach zu linearisieren, und bei diesem Durcheinanderschütteln geht die Dimension 2 der Ebene $\mathbb{R}^2$ verloren.

In der Tat gibt es keine stetigen (s. u.) bijektiven Abbildungen zwischen $\mathbb{R}^2$ und $\mathbb{R}$, und allgemeiner zwischen verschiedendimensionalen Kontinua. Dieser Satz, den Dedekind unmittelbar nach der brieflichen Mitteilung von $|\mathbb{I}^n| = |\mathbb{I}|$ durch Cantor vermutet hatte, wurde erst 1911 durch Luitzen Brouwer (1881 – 1966) vollständig bewiesen. Wir diskutieren am Ende des Kapitels noch eine stetige Surjektion von $\mathbb{R}$ nach $\mathbb{R}^2$, die allerdings nicht injektiv ist – und nicht sein kann.

Daß es etwa keine stetige Bijektion $f: \mathbb{I}^2 \rightarrow \mathbb{I}$ mit $\mathbb{I} = [0, 1] \subseteq \mathbb{R}$ geben kann, läßt sich noch relativ einfach zeigen. Für Leser, die einige Begriffe der Topologie kennen, sei hier der Beweis skizziert: Stetige Funktionen erhalten den Zusammenhang, und $\mathbb{I}^2$ ist nach Entfernung eines Punktes x mit $0 < f(x) < 1$ zusammenhängend, während $\mathbb{I}$ nach Entfernung von $f(x)$ zwei Komponenten hat. Also kann ein stetiges bijektives $f: \mathbb{I}^2 \rightarrow \mathbb{I}$ nicht existieren. (Der Beweis zeigt stärker, daß jedes stetige $f: \mathbb{I}^2 \rightarrow \mathbb{I}$ jeden Wert $x \in \text{rng}(f)$ überabzählbar oft annimmt mit Ausnahme von allenfalls zwei Werten.)

Es folgt, daß es dann auch kein stetiges bijektives $g: \mathbb{I} \rightarrow \mathbb{I}^2$ gibt, denn eine derartige Funktion g hätte automatisch eine stetige bijektive Umkehrabbildung $g^{-1}: \mathbb{I}^2 \rightarrow \mathbb{I}$. (Umkehrungen von stetigen Bijektionen brauchen nicht stetig zu sein. Sei etwa $g: [0, 2\pi[\rightarrow \mathbb{K}$ mit $g(\alpha) = (\cos(\alpha), \sin(\alpha))$, oder g irgendeine Aufwicklung eines halboffenen Intervalls zu einer Kreislinie; g^{-1} ist nicht stetig in $g(0)$. Ist der Definitionsbereich eines stetigen bijektiven g kompakt und der Zielraum Hausdorffsch, so hat g automatisch eine stetige Umkehrabbildung.)

Übung

Für alle natürlichen Zahlen n , $m \geq 1$ ist $|\mathbb{R}^n| = |\mathbb{R}^m|$.

Alternativ zu einem induktiven Beweis kann man die Idee der Verschmelzung zweier reeller Zahlen zu einer durch „Mischen“ oder „Einfädeln“ der Nachkommastellen verallgemeinern zu einer Verschmelzung von n reellen Zahlen zu einer – und sogar zu einer Verschmelzung von abzählbar vielen reellen Zahlen zu einer, wie wir gleich zeigen werden.

Bei der Umkehrung dieser Idee – aus einer reellen Zahl zwei zu machen – ist etwas Vorsicht geboten. Aus

$z = c, c_0 c_1 c_2 \dots$ können wir zwar

$x = a, c_0 c_2 c_4 \dots$ und $y = b, c_1 c_3 c_5 \dots$, wobei hier $(a, b) = g^{-1}(c)$ ist,

herauslösen. x und y sind aber nicht mehr notwendig in kanonischer Darstellung.

Sei etwa

$$g(0, 0) = 0, \quad g(1, 0) = 1,$$

$$z_0 = 1, 0 \ 1 \ 0 \ 1 \ 0 \ 1 \dots,$$

$$z_1 = 0, 9 \ 1 \ 9 \ 1 \ 9 \ 1 \dots$$

Dann ist

$$x_0 = 1, 0 \ 0 \ 0 \dots, \quad y_0 = 0, 1 \ 1 \ 1 \dots,$$

$$x_1 = 0, 9 \ 9 \ 9 \dots, \quad y_1 = 0, 1 \ 1 \ 1 \dots$$

Also $x_0 = x_1$ und $y_0 = y_1$. Aber $z_0 \neq z_1$! Also ist diese Teilungsfunktion nicht notwendig injektiv.

Die folgende Übung zeigt einen Weg, doch direkt eine Bijektion von $\mathbb{R}$ nach $\mathbb{R}^2$ durch Aufspaltung der Dezimaldarstellung einer reellen Zahl zu erhalten. Es ergibt sich ein Beweis des Satzes, der den Satz von Cantor-Bernstein nicht heranzieht. Die Idee stammt von Julius König, Cantor hat diesen Trick übersehen.

Übung (Trick von Julius König)

Ein *Block* einer reellen Zahl $x = b, a_0 a_1 a_2 \dots$ in kanonischer Darstellung ist eine endliche Folge $a_n, a_{n+1}, \dots, a_{n+m}$ aus Nachkommastellen mit der Eigenschaft:

$$a_{n-1} \neq 0 \text{ (falls } n > 0), \quad a_n = \dots = a_{n+m-1} = 0, \quad a_{n+m} \neq 0.$$

Beginnt z. B. die Dezimaldarstellung von x mit $1,100130710001 \dots$, so sind $1, 001, 3, 07, 1, 0001$ die ersten Blöcke von x .

Konstruieren Sie mit Hilfe von Blöcken eine Bijektion zwischen

I und $I \times I$, wobei $I = \{x \in \mathbb{R} \mid 0 < x \leq 1\}$.

[Aufspaltung der Blöcke anstelle der Ziffern.]

Julius König fand den Trick vor 1900. Die erste dem Autor bekannte Referenz ist [Schoenflies 1900, S. 23], wo es heißt: „... und zweitens denke man sich die eventuellen Nullen mit der ersten auf sie folgenden Ziffer [ungleich 0] zu je einer Gruppe verbunden, und dehne das Abbildungsgesetz [das Mischverfahren] auf diese Zahlengruppen aus¹⁾.“ Die Fußnote 1) hierzu ist: „1) Dieser Gedanke rührt von J. König her.“

Das Multiplikationsproblem

Eine gewagte, aber natürliche Frage an dieser Stelle ist nun :

Gilt für unendliche Mengen M immer $|M \times M| = |M|$?

Das triviale Argument aus obigem Beweis zeigt: $|M| \leq |M \times M|$. In der anderen Richtung haben wir Schwierigkeiten. Dennoch ist die Antwort auf die Frage „ja“, wie wir später zeigen werden, nachdem wir weitere Beispiele für diese Gleichung kennengelernt haben.

Folgen reeller Zahlen

Zunächst eine häufig gebrauchte Notation.

Definition (die Menge ${}^A B$)

Seien A, B Mengen. Wir setzen:

$${}^A B = \{ f \mid f : A \rightarrow B \}.$$

${}^{\mathbb{N}}\mathbb{R}$ ist also die Menge aller Funktionen von $\mathbb{N}$ nach $\mathbb{R}$ oder anders betrachtet, die Menge aller Folgen $x_0, x_1, \dots, x_n, \dots, n \in \mathbb{N}$, reeller Zahlen.

Übung

Seien A, B, C Mengen mit $|B| = |C|$.

Dann gilt $|{}^A B| = |{}^A C|$ und $|{}^B A| = |{}^C A|$.

Satz

Es gilt $|{}^{\mathbb{N}}\mathbb{R}| = |\mathbb{R}|$.

Beweis

Für $x \in \mathbb{R}$ sei $f_x \in {}^{\mathbb{N}}\mathbb{R}$ definiert durch: $f_x(n) = x$ für alle $n \in \mathbb{N}$.

Dann ist $F : \mathbb{R} \rightarrow {}^{\mathbb{N}}\mathbb{R}$ mit $F(x) = f_x$ injektiv. Also gilt $|\mathbb{R}| \leq |{}^{\mathbb{N}}\mathbb{R}|$.

Wir zeigen nun $|{}^{\mathbb{N}}\mathbb{R}| \leq |\mathbb{R}|$. Sei $I = [0, 1] = \{ x \in \mathbb{R} \mid 0 \leq x \leq 1 \}$.

Wegen $|I| = |\mathbb{R}|$ genügt es zu zeigen:

$$|{}^{\mathbb{N}}I| \leq |\mathbb{R}|.$$

Wir definieren $F : {}^{\mathbb{N}}I \rightarrow \mathbb{R}$. Sei hierzu $f \in {}^{\mathbb{N}}I$, also $f : \mathbb{N} \rightarrow I$.

Sei in kanonischen Dezimaldarstellungen:

$$f(0) = 0, a_{0,0} \ a_{0,1} \ a_{0,2} \ \dots$$

$$f(1) = 0, a_{1,0} \ a_{1,1} \ a_{1,2} \ \dots$$

$$f(2) = 0, a_{2,0} \ a_{2,1} \ a_{2,2} \ \dots$$

$$f(3) = 0, a_{3,0} \ a_{3,1} \ a_{3,2} \ \dots$$

$\dots$

$$f(n) = 0, a_{n,0} \ a_{n,1} \ a_{n,2} \ \dots$$

Wir definieren

$$F(f) = 0, a_{0,0} a_{0,1} a_{1,0} a_{0,2} a_{1,1} a_{2,0} a_{0,3} a_{1,2} a_{2,1} a_{3,0} a_{0,4} \dots,$$

d.h. die Nachkommastellen von $F(f)$ werden gegeben durch die Cantorsche Diagonalaufzählung $\pi : \mathbb{N}^2 \rightarrow \mathbb{N}$ der Nachkommastellen der $f(n)$. Genauer gilt, daß die $(n+1)$ -te Nachkommastelle von $F(f)$ definiert ist als $a_{k,\ell}$, wobei $\pi^{-1}(n) = (k, \ell)$.

$F(f)$ ist in kanonischer Darstellung für alle $f \in {}^{\mathbb{N}}\mathbb{I}$ (!).

- Offenbar ist also $F : {}^{\mathbb{N}}\mathbb{I} \rightarrow \mathbb{R}$ injektiv.

Selbst ein abzählbar unendlich dimensionales Kontinuum hat also die Größe von $\mathbb{R}$. Auch diesen Sachverhalt hat Cantor herausgestellt (vgl. den Briefauszug am Ende des Kapitels).

$\mathbb{R}$ und die Potenzmenge der natürlichen Zahlen

Wir zeigen $|\mathbb{R}| = |\mathcal{P}(\mathbb{N})|$. Die Grundidee ist, daß wir einer Teilmenge A von $\mathbb{N}$ ihre Indikatorfunktion $\text{ind}_A : \mathbb{N} \rightarrow \{0, 1\}$ zuordnen. Allgemein definieren wir:

Definition (*Indikatorfunktion oder charakteristische Funktion*)

Sei M eine Menge, und sei $A \subseteq M$. Dann ist die

Indikatorfunktion von A in M $\text{ind}_{A,M} : M \rightarrow \{0, 1\}$ definiert durch

$$\text{ind}_{A,M}(x) = \begin{cases} 1, & \text{falls } x \in A, \\ 0, & \text{falls } x \notin A. \end{cases}$$

Es gilt nun:

Satz

Sei M eine Menge. Dann gilt $|\mathcal{P}(M)| = |^M\{0, 1\}|$.

Beweis

- Wir definieren $f : \mathcal{P}(M) \rightarrow {}^M\{0, 1\}$ bijektiv durch $f(A) = \text{ind}_{A,M}$ für $A \subseteq M$.

Hinsichtlich $|\mathcal{P}(\mathbb{N})| = |\mathbb{R}|$ fassen wir nun $\text{ind}_{A,\mathbb{N}}$ für $A \subseteq \mathbb{N}$ einfach als reelle Zahl im Einheitsintervall in Binärdarstellung auf. Da endliche Mengen dadurch in eine trivial endende Darstellung einer reellen Zahl übergehen, ist aber etwas Vorsicht geboten. Wir brauchen eine Vorüberlegung.

Definition ($\mathcal{P}^*(\mathbb{N})$)

Wir setzen

$$\mathcal{P}^*(\mathbb{N}) = \{A \subseteq \mathbb{N} \mid A \text{ ist unendlich}\}.$$

Als Übung kann der Leser versuchen, folgenden Satz zu zeigen:

Satz

Es gilt $|\mathcal{P}^*(\mathbb{N})| = |\mathcal{P}(\mathbb{N})|$.

Beweis

$|\mathcal{P}^*(\mathbb{N})| \leq |\mathcal{P}(\mathbb{N})|$ ist klar wegen $\mathcal{P}^*(\mathbb{N}) \subseteq \mathcal{P}(\mathbb{N})$.

Wir zeigen $|\mathcal{P}(\mathbb{N})| \leq |\mathcal{P}^*(\mathbb{N})|$.

Hierzu definieren wir $f: \mathcal{P}(\mathbb{N}) \rightarrow \mathcal{P}^*(\mathbb{N})$ wie folgt.

Sei $U = \{2n + 1 \mid n \in \mathbb{N}\}$ die Menge aller ungeraden natürlichen Zahlen.

Wir setzen für $A \subseteq \mathbb{N}$:

$$f(A) = \{2n \mid n \in A\} \cup U.$$

- Dann ist $f: \mathcal{P}(\mathbb{N}) \rightarrow \mathcal{P}^*(\mathbb{N})$ injektiv, also gilt $|\mathcal{P}(\mathbb{N})| \leq |\mathcal{P}^*(\mathbb{N})|$.

Der Beweis, den der Leser gefunden hat, ist vielleicht: $\mathcal{P}(\mathbb{N})$ ist überabzählbar (durch Diagonalisierung, vgl. auch Kapitel 10), $E = \{A \subseteq \mathbb{N} \mid A \text{ ist endlich}\}$ ist abzählbar, und die Subtraktion einer abzählbaren Menge ändert die Mächtigkeit nicht.

Damit können wir nun leicht den fundamentalen Zusammenhang zwischen den Mächtigkeiten der natürlichen und der reellen Zahlen zeigen:

Satz

Es gilt $|\mathbb{R}| = |\mathcal{P}(\mathbb{N})|$.

Beweis

Sei $I =]0, 1] = \{x \in \mathbb{R} \mid 0 < x \leq 1\}$.

Es gilt $|I| = |\mathbb{R}|$ und $|\mathcal{P}(\mathbb{N})| = |\mathcal{P}^*(\mathbb{N})|$.

Es genügt also zu zeigen, daß $|I| = |\mathcal{P}^*(\mathbb{N})|$.

Hierzu definieren wir $f: I \rightarrow \mathcal{P}^*(\mathbb{N})$ wie folgt.

Sei $x \in I$ und sei, in kanonischer Binärdarstellung,

$$x = 0, a_0 a_1 a_2 \dots$$

mit $a_n \in \{0, 1\}$. Wir definieren $f(x) \subseteq \mathbb{N}$ durch:

$$f(x) = \{n \mid a_n = 1\}.$$

Da die kanonische Darstellung von x nicht trivial endet,

ist $f(x)$ unendlich für alle $x \in I$.

- Das so definierte $f: I \rightarrow \mathcal{P}^*(\mathbb{N})$ ist bijektiv.

Die beiden wichtigsten Strukturen der Mathematik $\mathbb{N}$ und $\mathbb{R}$ sind also von unterschiedlicher Mächtigkeit, und die Mächtigkeit der zweiten ist gerade die Mächtigkeit der Potenzmenge der ersten. Ein bemerkenswerter Zusammenhang! In der Mengenlehre wird oftmals sogar eine Teilmenge von $\mathbb{N}$ direkt als reelle Zahl bezeichnet.

Eine etwas andere Strategie, um $|\mathbb{R}| = |\mathcal{P}(\mathbb{N})|$ zu zeigen, ist diese: Der unproblematische Teil ist der Nachweis von $|\mathbb{R}| = |[0, 1]| \leq |\mathcal{P}(\mathbb{N})|$, was man durch kanonische binäre Darstellung von $x \in [0, 1]$ leicht zeigt. Für die Umkehrung $|\mathcal{P}(\mathbb{N})| \leq |\mathbb{R}|$ ist die Mehrdeutigkeit der Binärdarstellung hinderlich. Die-

ses Problem kann man nun umgehen, indem man zu b -adischen Entwicklungen mit $b > 2$ übergeht. Der einfachste Fall $b = 3$ führt zur sogenannten *Cantormenge*, die wir in Kapitel 12 des zweiten Abschnitts ausführlich untersuchen werden.

Übung

Die *Cantormenge* $C \subseteq \mathbb{R}$ ist definiert als die Menge aller reellen Zahlen x mit $0 \leq x \leq 1$, für die es eine (nicht notwendig kanonische) Ternär-Darstellung (= 3-adische Darstellung)

$$x = 0, a_0 a_1 a_2 \dots$$

gibt mit der Eigenschaft: für alle n ist $a_n \neq 1$.

Anders ausgedrückt: x läßt sich schreiben als

$$x = a_0/3 + a_1/3^2 + a_2/3^3 + \dots$$

mit $a_n \in \{0, 2\}$.

Man zeige: $|C| = |\mathbb{N}\{0, 1\}| = |\mathcal{P}(\mathbb{N})|$ (und damit $|\mathcal{P}(\mathbb{N})| = |\mathbb{R}|$).

Die Gleichung $|\mathcal{P}(M) \times \mathcal{P}(M)| = |\mathcal{P}(M)|$

Wir erhalten aus $|\mathbb{R}| = |\mathcal{P}(\mathbb{N})|$ auch einen neuen Beweis von $|\mathbb{R} \times \mathbb{R}| = |\mathbb{R}|$. Nützlich hierfür ist eine einfache Definition.

Definition (2M)

Sei M eine Menge. Wir setzen:

$$2M = M \times \{0\} \cup M \times \{1\}.$$

$2M = M \times \{0, 1\}$ ist also die Vereinigung zweier disjunkter „Kopien“ von M . Es gilt nun:

Satz

Sei M eine Menge und es gelte $|2M| = |M|$.

Dann gilt $|\mathcal{P}(M) \times \mathcal{P}(M)| = |\mathcal{P}(M)|$.

Beweis

Sei $f: 2M \rightarrow M$ bijektiv.

Definiere $g: \mathcal{P}(M) \times \mathcal{P}(M) \rightarrow \mathcal{P}(M)$ durch

$$g(A, B) = f''(A \times \{0\} \cup B \times \{1\})$$

für $A, B \subseteq M$.

– Dann ist $g: \mathcal{P}(M) \times \mathcal{P}(M) \rightarrow \mathcal{P}(M)$ bijektiv.

Die Voraussetzung $|2M| = |M|$ ist erfüllt, falls $|M \times M| = |M|$ gilt und M mehr als ein Element hat. Die Eigenschaft $|M \times M| = |M|$ vererbt sich also von einer unendlichen Menge auf ihre Potenzmenge.

Sehr leicht folgt nun:

neuer Beweis von $|\mathbb{R} \times \mathbb{R}| = |\mathbb{R}|$

Es gilt $|2\mathbb{N}| = |\mathbb{N}|$. Also $|\mathcal{P}(\mathbb{N}) \times \mathcal{P}(\mathbb{N})| = |\mathcal{P}(\mathbb{N})|$.

– Wegen $|\mathbb{R}| = |\mathcal{P}(\mathbb{N})|$ folgt die Behauptung.

Eine einfache Verallgemeinerung liefert auch das Resultat $|\mathbb{N}\mathbb{R}| = |\mathbb{R}|$:

Übung

Sei M eine Menge. Zeigen Sie:

(i) $|M \times \mathbb{N}| = |M|$ folgt $|\mathbb{N}\mathcal{P}(M)| = |\mathcal{P}(M)|$.

[Analog zu: $|2M| = |M|$ folgt $|\mathcal{P}(M) \times \mathcal{P}(M)| = |\mathcal{P}(M)|$.]

(ii) Folgern Sie hiermit $|\mathbb{N}\mathbb{R}| = |\mathbb{R}|$.

Schließlich halten wir fest:

Satz (Die Mächtigkeit der Folgen in $\mathbb{N}$)

Es gilt $|\mathbb{N}\mathbb{N}| = |\mathbb{R}|$.

Beweis

Wir haben nach den bisherigen Resultaten und wegen $\mathbb{N}\mathbb{N} \subseteq \mathbb{N}\mathbb{R}$:

$$|\mathbb{R}| = |\mathcal{P}(\mathbb{N})| = |\mathbb{N}\{0, 1\}| \leq |\mathbb{N}\mathbb{N}| \leq |\mathbb{N}\mathbb{R}| = |\mathbb{R}|.$$

– Also folgt die Behauptung nach Cantor-Bernstein.

$\mathbb{N}\mathbb{N}$ ist die Menge aller Folgen $a_0, a_1, a_2, \dots, a_n, \dots$ von natürlichen Zahlen. Es folgt, daß $|\mathbb{N}\mathbb{N}| = |\mathbb{N}\mathbb{R}| = |\mathcal{P}(\mathbb{N})|$, es gibt also nicht mehr Folgen natürlicher oder sogar reeller Zahlen als Teilmengen von $\mathbb{N}$.

Baireraum und Cantorraum

Definition (Baireraum und Cantorraum)

$\mathbb{N}\mathbb{N}$ heißt der *Baireraum*, $\mathbb{N}\{0, 1\}$ der *Cantorraum*.

Der Baireraum – benannt nach René Baire (1874 – 1932) – und der Cantorraum sind in der Mengenlehre von großer Bedeutung. In vielen Untersuchungen ersetzen sie die reellen Zahlen. Die Zuordnung einer reellen Zahl $x \in I = [0, 1]$ zur Folge $b \in \mathbb{N}\{0, 1\}$ ihrer Nachkommastellen in Binärdarstellung ist nicht immer eindeutig. In der Analysis ist $\mathbb{R}$ als stetige Linie fundamental, in der Mengenlehre ist das Phänomen der Uneindeutigkeit eher lästig. Die Folgenräume haben aber ganz ähnliche Eigenarten wie die reellen Zahlen: Wie eine reelle Zahl durch Angabe von immer mehr Nachkommastellen immer genauer beschrieben wird, so werden Elemente f der Folgenräume durch Angabe von immer längeren Anfangsstücken $f(0), f(1), \dots, f(n)$ immer besser approximiert. Der Leser kann sich die Elemente der Folgenräume als Information vorstellen, die portionsweise und insgesamt abzählbar oft gesammelt wird. Bei Elementen des

Sei $B = \{g \in {}^{\mathbb{N}}\{0, 1\} \mid \text{es gibt ein } n_0 \in \mathbb{N} \text{ mit:}$
 $\quad g(n) = 1 \text{ für alle } n \geq n_0 \text{ oder } g(n) = 0 \text{ für alle } n \geq n_0\}.$

Dann ist $F^* : {}^{\mathbb{N}}\mathbb{N} \rightarrow {}^{\mathbb{N}}\{0, 1\} - B$ bijektiv. Wie oben sei

$H^*(g) = 0, g(0)g(1)\dots$ in Binärdarstellung.

für $g \in {}^{\mathbb{N}}\{0, 1\} - B$. Dann ist $H^* : {}^{\mathbb{N}}\{0, 1\} - B \rightarrow [0, 1[- C$ bijektiv, wobei hier $C = \{x \in [0, 1[\mid \text{es gibt eine endliche Binärdarstellung von } x\}$. H^* erhält nun zudem die Nähebeziehungen in perfekter Weise, wie sich der Leser leicht überlegt.

Insgesamt zeigen die Überlegungen, daß der Baireraum zu den reellen Zahlen im Einheitsintervall, die keine endliche Binärdarstellung besitzen, strukturell äquivalent ist. Mit Hilfe von Kettenbrüchen kann man ein gleichwertiges Ergebnis erzielen: Einem Element f des Baireraums wird durch einen Kettenbruch die Zahl $x = K(f)$ wie im Diagramm zugeordnet, wobei $f'(n) = f(n) + 1$ für alle $n \in \mathbb{N}$. In der Analysis zeigt man, daß jedes solche $K(f)$ eine irrationale Zahl ist, daß jede irrationale Zahl x des Einheitsintervalls als Kettenbruch geschrieben werden kann, und daß die Nähebeziehungen durch die Kettenbruchzuordnung perfekt erhalten bleiben. Der Baireraum ist damit strukturell äquivalent zu den irrationalen Zahlen des Einheitsintervalls.

$$x = \frac{1}{f'(0) + \frac{1}{f'(1) + \frac{1}{f'(2) + \dots}}}$$

Die Mächtigkeit der reellen Funktionen

Wir haben $|\mathbb{N}| < |\mathbb{R}|$ gezeigt. Die natürliche Frage ist nun:

Gibt es Mengen M mit $|\mathbb{R}| < |M|$?

Die Antwort ist *ja*. Cantors Diagonalargument kann man verwenden, um zu zeigen, daß die Menge der reellen Funktionen größer ist als $\mathbb{R}$:

Definition (reelle Funktionen)

Eine *reelle Funktion* ist ein $f : \mathbb{R} \rightarrow \mathbb{R}$.

Die Menge aller reellen Funktionen bezeichnen wir mit $\mathfrak{F}$.

Es gilt also $\mathfrak{F} = {}^{\mathbb{R}}\mathbb{R}$. Wir zeigen nun:

Satz (über die Mächtigkeit der reellen Funktionen)

Es gilt $|\mathbb{R}| < |\mathfrak{F}|$.

Beweis

Zunächst gilt $|\mathbb{R}| \leq |\mathfrak{F}|$: Wir setzen für alle $x \in \mathbb{R}$ $g(x) = f_x \in \mathfrak{F}$, wobei $f_x(y) = x$ für alle $y \in \mathbb{R}$. Dann ist $g : \mathbb{R} \rightarrow \mathfrak{F}$ injektiv.

Sei nun $F : \mathbb{R} \rightarrow \mathfrak{F}$ beliebig. Wir zeigen, daß F nicht surjektiv ist. Wir definieren hierzu eine reelle Funktion d wie folgt. Für $x \in \mathbb{R}$ sei

$$d(x) = F(x)(x) + 1.$$

[Es gilt $F(x) \in \mathfrak{F}$ für alle $x \in \mathbb{R}$. $F(x)$ ist also eine reelle Funktion, die wir an der Stelle x auswerten können. Der um eins erhöhte Wert dieser Auswertung ist $d(x)$.]

Wir zeigen, daß $d : \mathbb{R} \rightarrow \mathbb{R}$ nicht im Wertebereich der Funktion F liegt.

Annahme doch. Sei also $y \in \mathbb{R}$ mit $F(y) = d$. Dann gilt

$$F(y)(x) = d(x) \text{ für alle } x \in \mathbb{R}.$$

Insbesondere gilt $F(y)(y) = d(y)$.

Aber $d(y) = F(y)(y) + 1$, *Widerspruch!*

– Also ist $d \notin \text{rng}(F)$, und damit F nicht surjektiv.

Aus der Überabzählbarkeit der reellen Zahlen und der Abzählbarkeit der algebraischen Zahlen haben wir die Existenz von transzendenten Zahlen gewonnen. Nun haben wir „über-reell“ für die Größe der Menge der reellen Funktionen bewiesen. Kann man ein Analogon finden zum Beweis der Existenz von transzendenten Zahlen?

In gewisser Weise ist das möglich. Für diese Ausführungen müssen wir allerdings beim Leser einige Kenntnisse der reellen Analysis voraussetzen.

Wir betrachten zunächst die Größe bestimmter natürlicher Teilmengen von $\mathfrak{F}$. Klar ist, daß die Menge aller konstanten Funktionen, d.h. aller $f_c : \mathbb{R} \rightarrow \mathbb{R}$ mit $f_c(x) = c$ für alle $x \in \mathbb{R}$ für ein gewisses $c \in \mathbb{R}$, die Größe von $\mathbb{R}$ hat.

Komplizierter ist schon die Menge $\mathfrak{S} = \{f \in \mathfrak{F} \mid f \text{ ist stetig}\}$ der stetigen reellen Funktionen. Intuitiv bedeutet die Stetigkeit einer reellen Funktion f im Punkt a , daß $f(x)$ nahe bei $f(a)$ liegt, wenn x nahe bei a ist. Die genaue Definition ist:

Ein $f : \mathbb{R} \rightarrow \mathbb{R}$ heißt *stetig in einem Punkt* $a \in \mathbb{R}$, wenn gilt:

Für alle $\varepsilon > 0$ existiert ein $\delta > 0$, sodaß für alle x gilt:

$$|x - a| < \delta \text{ folgt } |f(x) - f(a)| < \varepsilon.$$

f heißt *stetig*, falls f stetig in allen $a \in \mathbb{R}$ ist.

Aus dieser Bedingung folgt nun aber, daß eine stetige Funktion bereits durch ihre Werte auf $\mathbb{Q}$ eindeutig bestimmt ist:

Übung

Sind $f, g \in \mathfrak{S}$, und ist $f|_{\mathbb{Q}} = g|_{\mathbb{Q}}$, so gilt $f = g$.

Dies bedeutet, daß es höchstens so viele stetige Funktionen gibt wie Funktionen $f : \mathbb{Q} \rightarrow \mathbb{R}$. Nach unseren Ergebnissen aus dem letzten Kapitel ist aber

$$|\{f \mid f : \mathbb{Q} \rightarrow \mathbb{R}\}| = |\mathbb{Q}^{\mathbb{R}}| = |\mathbb{N}^{\mathbb{R}}| = |\mathbb{R}|.$$

Also gilt $|\mathfrak{S}| \leq |\mathbb{R}|$. Andererseits ist jede konstante Funktion auf $\mathbb{R}$ stetig, also gilt $|\mathbb{R}| \leq |\mathfrak{S}| \leq |\mathbb{R}|$, und damit $|\mathfrak{S}| = |\mathbb{R}|$. Es gibt also lediglich so viele stetige Funktionen wie reelle Zahlen. Da jede differenzierbare Funktion stetig und jede

konstante Funktion differenzierbar ist, folgt auch, daß die differenzierbaren Funktionen die Mächtigkeit von $\mathbb{R}$ haben. [Differenzierbare Funktionen sind anschaulich stetige Funktionen ohne „Knicke“.]

Hessenberg (1906, § 29): „... In analoger Weise läßt sich beweisen, daß die Menge aller stetigen Funktionen von der Mächtigkeit des Kontinuums ist. Überraschend sind diese Resultate aus dem gleichen Grunde wie die Abzählbarkeit der rationalen Zahlen, weil offenbar die Anordnung der Punkte eines Raumes durch die Zuordnung in das Kontinuum völlig zerstört wird, während umgekehrt die Menge der stetigen Funktionen eine Ordnung erhält, die ihr nach der ursprünglichen Definition nicht zukommt.

Läßt man die Beschränkung der Stetigkeit fallen und betrachtet die Menge aller [reellen] Funktionen ..., so ist diese ... von größerer Mächtigkeit als das Kontinuum.

Hiermit sind drei Mengen aufgewiesen, die schon lange vor Schöpfung der Mengenlehre Gegenstand mathematischer Arbeit waren: die Menge der ganzen Zahlen, der reellen Zahlen und der Funktionen. Sie sind nicht erst zu dem Zweck konstruiert, die Möglichkeit verschiedener Mächtigkeiten darzutun, vielmehr boten sie sogleich der Mengenlehre einen fruchtbaren Anknüpfungspunkt an vorhandene Arbeitsgebiete.“

Nun kann man aber überraschenderweise zeigen, daß die Menge der Riemann-integrierbaren Funktionen die Mächtigkeit von $\mathfrak{F}$ besitzt.

Übung (Voraussetzung: Kenntnis des Begriffs „Riemann-integrierbar“)

Beweisen Sie diese Behauptung.

[Betrachten Sie die Cantormenge $C \subseteq [0, 1]$ und Funktionen

$f: [0, 1] \rightarrow [0, 1]$, die außerhalb von C gleich 0 sind. Es gilt $|C| = |\mathbb{R}|$.]

Dagegen ist die Menge aller $f \in \mathfrak{F}$, die durch eine abzählbare Menge von stetigen Funktionen eindeutig beschreibbar/approximierbar sind, von der Mächtigkeit $|\mathbb{N}\mathbb{R}| = |\mathbb{R}|$. Dies zeigt den „transzendenten“ Charakter der integrierbaren Funktionen: allein ihre Anzahl bringt schon mit sich, daß es integrierbare Funktionen f gibt, die nicht durch eine Folge $f_0, f_1, \dots, f_n, \dots$ von stetigen Funktionen punktweise approximiert werden können. Das gleiche gilt für jedes Reservoir von approximierenden Funktionen der Größe $\mathbb{R}$.

Wie gelangt man zu einer Menge größerer Mächtigkeit?

Wir haben bisher $|\mathbb{N}| < |\mathbb{R}|$ und $|\mathbb{R}| < |\mathfrak{F}|$ gezeigt.

Gibt es ein allgemeines Prinzip oder eine Operation, um von einer beliebigen Menge M zu einer Menge $\mathcal{M}$ mit größerer Mächtigkeit als M zu gelangen?

$\mathcal{M} = M \times M$ ist ungeeignet, wie wir für $M = \mathbb{N}$ und $M = \mathbb{R}$ gesehen haben.

Jedoch gilt:

$$|\mathbb{R}| = |\mathcal{P}(\mathbb{N})|.$$

Wie sieht es nun mit dem Verhältnis von $\mathbb{R}$ und $\mathfrak{F}$ aus? In der Tat gilt hier eine zu $\mathbb{N}$ und $\mathbb{R}$ analoge Beziehung:

Satz

Es gilt $|\mathfrak{F}| = |\mathcal{P}(\mathbb{R})|$.

Beweis

Jedes $f \in \mathfrak{F}$ ist eine Teilmenge von $\mathbb{R} \times \mathbb{R}$, also gilt $\mathfrak{F} \subseteq \mathcal{P}(\mathbb{R} \times \mathbb{R})$, also

$$|\mathfrak{F}| \leq |\mathcal{P}(\mathbb{R} \times \mathbb{R})| = |\mathcal{P}(\mathbb{R})|,$$

wegen $|\mathbb{R} \times \mathbb{R}| = |\mathbb{R}|$.

Andererseits können wir jedem $A \subseteq \mathbb{R}$ die Funktion $F(A) = \text{ind}_{A, \mathbb{R}} \in \mathfrak{F}$ zuordnen, d.h. es gilt für $x \in \mathbb{R}$:

$$F(A)(x) = \begin{cases} 1, & \text{falls } x \in A, \\ 0, & \text{falls } x \notin A. \end{cases}$$

Offenbar ist dann $F : \mathcal{P}(\mathbb{R}) \rightarrow \mathfrak{F}$ injektiv, also $|\mathcal{P}(\mathbb{R})| \leq |\mathfrak{F}|$.

– Insgesamt also $|\mathfrak{F}| = |\mathcal{P}(\mathbb{R})|$.

Übung

Zeigen Sie

$$(i) \quad |\mathfrak{F} \times \mathfrak{F}| = |\mathfrak{F}|,$$

$$(ii) \quad |\mathbb{R}^{\mathfrak{F}}| = |\mathfrak{F}|.$$

Wir haben $|\mathbb{N}| < |\mathcal{P}(\mathbb{N})|$ und $|\mathbb{R}| < |\mathcal{P}(\mathbb{R})|$. Die Potenzmengenoperation ist also ein guter Kandidat für ein allgemeines Prinzip zur Erzeugung von größeren Mächtigkeiten. Bereits im Endlichen liefert sie exponentielles Wachstum: Es gilt $|\mathcal{P}(\bar{n})| = |\bar{n}|^{\bar{n}}$ für alle $n \in \mathbb{N}$, wobei hier wieder $\bar{n} = \{0, 1, \dots, n-1\}$. Die Potenzmenge einer Menge mit n Elementen hat also 2^n Elemente.

Wir beschäftigen uns mit der Potenzmengenoperation im nächsten Kapitel genauer.

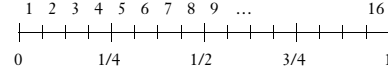
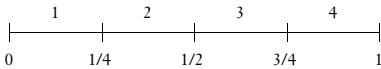
Anhang: Eine stetige Surjektion von $[0, 1]$ nach $[0, 1] \times [0, 1]$

Sei $I = [0, 1] = \{x \in \mathbb{R} \mid 0 \leq x \leq 1\}$ das reelle Einheitsintervall. Wir skizzieren hier für mit der Analysis ein wenig vertraute Leser die Konstruktion einer stetigen surjektiven Funktion $f : I \rightarrow I \times I$ (Details als Übung). Die erste derartige Funktion wurde von Giuseppe Peano 1890 gefunden. Die folgende Konstruktion geht auf David Hilbert (1891) zurück.

Zur Definition von f zerteilen wir zunächst iteriert I und $I \times I$ in je vier *abgeschlossene* Teilintervalle bzw. Teilquadrate, wobei wir die im n -ten Schritt entstandenen 4^n Teile einander bijektiv zuordnen. Die folgende Skizze zeigt die ersten drei Zerlegungen und die entsprechenden Zuordnungen:

1	2
4	3

1	4	5	6
2	3	8	7
15	14	9	10
16	13	12	11



Sei $x \in I$ und $n \in \mathbb{N}$. Dann gibt es ein oder zwei Teilintervalle der n -ten Zerlegung von I (in 4^n Teile), in denen x liegt. (Zwei solche Intervalle existieren genau dann, wenn $x = m/4^n$ für ein $0 < m < 4^n$ gilt). Sei $k(x, n)$ die kleinste Nummer eines x enthaltenden Teilintervalls, d.h.

$$k(x, n) = \min \{ m \geq 1 \mid x \in [(m-1)/4^n, m/4^n] \},$$

und sei $Q(k(x, n))$ das zugeordnete Teilquadrat von $I \times I$ (incl. Rand).

Wir setzen für $x \in I$

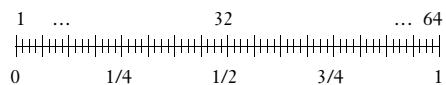
$$f(x) = \bigcap_{n \in \mathbb{N}} Q(k(x, n))$$

Dann gilt:

- (i) $f : I \rightarrow I \times I$ ist surjektiv,
- (ii) f ist stetig,
- (iii) f ist nicht injektiv.

[z. B. $f(1/6) = f(1/2) = (1/2, 1/2)$; geometrische Reihen sind hier nützlich.]

1	2	15	16	17			
4	3	14	13	18	...		
5	8	9	12				
6	7	10	11				



Georg Cantor über die Gleichung $|\mathbb{R}^2| = |\mathbb{R}|$, Brief an Dedekind vom 25. 6. 1877

„... Seit mehreren Jahren habe ich mit Interesse die Bemühungen verfolgt, denen man sich im Anschluß an Gauß, Riemann, Helmholtz und andern zur Klarstellung aller derjenigen Fragen hingegeben hat, welche die ersten Voraussetzungen der Geometrie betreffen. Dabei fiel mir auf, daß alle in dieses Feld schlagenden Untersuchungen ihrerseits von einer unbewiesenen Voraussetzung ausgehen, die mir nicht als selbstverständlich, vielmehr einer Begründung bedürftig erschienen ist. Ich meine die Voraussetzung, daß eine p -fach ausgedehnte stetige Mannigfaltigkeit zur Bestimmung ihrer Elemente p voneinander unabhängiger reeller Koordinaten bedarf; daß diese Zahl der Koordinaten für eine und dieselbe Mannigfaltigkeit weder vergrößert noch verkleinert werden könne.

Diese Voraussetzung war auch bei mir zu einer Ansichtssache geworden, ich war von ihrer Richtigkeit fast überzeugt; mein Standpunkt unterschied sich nur von allen anderen dadurch, daß ich jene Voraussetzung als einen Satz ansah, der eines Beweises in hohem Grade bedurfte und ich spitzte meinen Standpunkt zu einer Frage zu, die ich einigen Fachgenossen, im Besonderen auch bei Gelegenheit des Gaußjubiläums in Göttingen vorgelegt habe, nämlich zu folgender Frage: ‚Läßt sich ein stetiges Gebilde von p Dimensionen, wo $p > 1$, auf ein stetiges Gebilde von einer Dimension eindeutig beziehen, so daß jedem Punkte des einen ein und nur ein Punkt des anderen entspricht?‘

Die meisten, welchen ich diese Frage vorgelegt, wunderten sich sehr darüber, daß ich sie habe stellen können, da es sich ja von selbst verstünde, daß zur Bestimmung eines Punktes in einer Ausgedehntheit von p Dimensionen immer p unabhängige Koordinaten gebraucht werden. Wer jedoch in den Sinn der Frage eindrang mußte bekennen, daß es mindestens eines Beweises bedürfe, warum sie mit dem ‚selbstverständlichen‘ nein zu beantworten sei. Wie gesagt gehörte ich selbst zu denen, welche es für das Wahrscheinlichste hielten, daß jene Frage mit einem Nein zu beantworten sei, – bis ich vor ganz kurzer Zeit durch ziemlich verwickelte Gedankenreihen zu der Überzeugung gelangte, daß jene Frage ohne alle Einschränkung zu bejahen ist. Bald darauf fand ich den Beweis, welchen Sie heute vor sich sehen.

Da sieht man, welch’ wunderbare Kraft in den gewöhnlichen reellen rationalen und irrationalen Zahlen doch liegt, daß man durch sie im Stande ist die Elemente einer p -fach ausgedehnten stetigen Mannigfaltigkeit eindeutig mit einer einzigen Koordinate zu bestimmen; ja ich will nur gleich hinzufügen, daß ihre Kraft noch weiter geht, indem, wie Ihnen nicht entgehen wird, mein Beweis sich ohne besondere Vergrößerung der Schwierigkeiten auf Mannigfaltigkeiten mit einer unendlich großen Dimensionszahl ausdehnen läßt, vorausgesetzt, daß ihre unendlich vielen Dimensionen die Form einer einfach unendlichen Reihe bilden.

Nun scheint es mir, daß alle philosophischen oder mathematischen Deduktionen, welche von jener irrtümlichen Voraussetzung Gebrauch machen, unzulässig sind. Vielmehr wird der Unterschied, welcher zwischen Gebilden von verschiedener Dimensionszahl liegt, in ganz anderen Momenten gesucht werden müssen, als in der für charakteristisch gehaltenen Zahl der unabhängigen Koordinaten ...“

(Georg Cantor, Briefe (1991))

10. Die Mächtigkeit der Potenzmenge

Die Überlegungen zu Ende des letzten Kapitels führten uns zur folgenden Vermutung:

Ist M eine Menge, so ist die Potenzmenge von M von größerer Mächtigkeit als M .

Diese Vermutung ist nun in der Tat für alle Mengen richtig, nicht nur für $\mathbb{N}$ und $\mathbb{R}$. Der Leser kann sich auf einen kurzen und transparenten diagonalen Beweis dieser fundamentalen Tatsache freuen.

Der Satz von Cantor

Satz (Satz von Cantor über die Potenzmengenoperation)

Sei M eine Menge, $\mathcal{P}(M) = \{ X \mid X \subseteq M \}$ die Potenzmenge von M .

Dann gilt $|M| < |\mathcal{P}(M)|$.

Beweis

Offenbar ist $|M| \leq |\mathcal{P}(M)|$.

(Betrachte $F : M \rightarrow \mathcal{P}(M)$, definiert durch $F(x) = \{ x \}$ für $x \in M$.)

Sei nun $f : M \rightarrow \mathcal{P}(M)$ beliebig. Es genügt zu zeigen: f ist nicht surjektiv.

Wir setzen:

$D = \{ x \in M \mid x \notin f(x) \}$.

Dann ist $D \in \mathcal{P}(M)$. *Annahme*, $D \in \text{rng}(f)$.

Sei also $y \in M$ mit $f(y) = D$. Dann gilt:

$y \in D$ gdw $y \notin f(y)$ gdw $y \notin D$,

ersteres nach Definition von D , letzteres wegen $f(y) = D$.

– *Widerspruch!*

Wegen $|\mathbb{R}| = |\mathcal{P}(\mathbb{N})|$ und $|\mathbb{R}| = |\mathcal{P}(\mathbb{R})|$ liefert der Satz von Cantor auch einen neuen Beweis für die Überabzählbarkeit von $\mathbb{R}$ und für $|\mathbb{R}| < |\mathbb{R}|$.

Im zweiten Teil des Beweises wird $\text{rng}(f) \subseteq \mathcal{P}(M)$ nicht gebraucht. Der Beweis zeigt allgemein, daß wir für jede Menge M und jede Funktion f auf M eine Menge $D \subseteq M$ definieren können, die nicht im Wertebereich von f liegt:

Korollar (Lücken im Wertebereich)

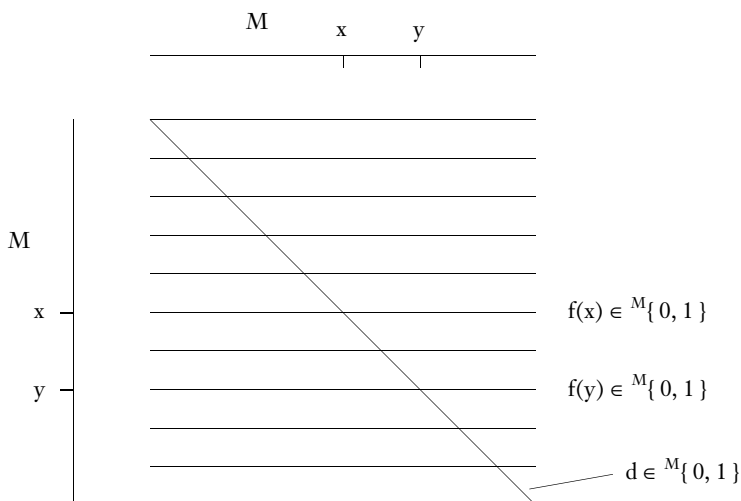
Sei M eine Menge, und sei f eine Funktion mit $\text{dom}(f) = M$.

Dann gilt $\{ x \in M \mid x \notin f(x) \} \notin \text{rng}(f)$.

Wie kommt man auf die Menge $D = \{x \in M \mid x \notin f(x)\}$? Bei genauerem Hinsehen erweist sich die Konstruktion von D als eine Diagonalisierung, wie sie uns in den Beweisen der Überabzählbarkeit von $\mathbb{R}$ und von $|\mathbb{R}| < |\mathfrak{F}|$ bereits begegnet ist: Wir identifizieren eine Teilmenge A von M mit ihrer Indikatorfunktion $\text{ind}_{A,M} : M \rightarrow \{0, 1\}$, wobei wieder $\text{ind}_{A,M}(x) = 1$ gdw $x \in A$. Die Potenzmenge von M wird dann zu ${}^M\{0, 1\}$, der Menge aller Indikatorfunktionen auf M . Sei nun $f : M \rightarrow {}^M\{0, 1\}$. Wir suchen ein $d \in {}^M\{0, 1\}$ mit $f(x) \neq d$ für alle $x \in M$. Wir können aber d verschieden von allen $f(x)$ konstruieren durch:

$$d(x) = \begin{cases} 1, & \text{falls } f(x)(x) = 0, \\ 0, & \text{falls } f(x)(x) = 1, \end{cases}$$

für alle $x \in M$. Dann gilt $d(x) \neq f(x)(x)$ für alle $x \in M$, also ist $d \notin \text{rng}(f)$.



Die Senkrechte des Diagramms repräsentiert M . Die Waagrechten seitlich der Senkrechten stehen für Funktionen $f(x) \in {}^M\{0, 1\}$, die man sich als 0-1-Folgen vorstellen kann. Die oberste Waagrechte ist der Definitionsbereich dieser Funktionen.

Die Diagonale steht für die konstruierte Funktion $d \in {}^M\{0, 1\}$ – ebenfalls eine 0-1-Folge. d ist in jedem $x \in M$ verschieden von $f(x)$, d. h. es gilt $f(x)(x) \neq d(x)$. $f(x)(x)$ ist der Wert der 0-1-Folge $f(x)$ an der Stelle x , d. h. der Wert der Waagrechten $f(x)$ an ihrem Schnittpunkt mit d . d ist dort gerade verschieden von diesem Wert, also ist d sicher nicht gleich $f(x)$. Und dies gilt für alle $x \in M$.

Übung

Sei $M = \{0, 1, 2, 3\}$. Bestimmen Sie $D \subseteq M$ wie im obigem Beweis für die Funktion $f : M \rightarrow \mathcal{P}(M)$ mit $f(0) = \{1, 3\}$, $f(1) = \{0, 2\}$, $f(2) = \{1, 2\}$, $f(3) = \{0, 1, 2\}$. Zeichnen Sie zudem obiges Diagramm für diese Situation mit 0-1-Folgen für $f(x)$ und bestimmen Sie d .

Durch iterierte Anwendung der Potenzmengenoperation können wir nun, ausgehend von einer beliebigen Menge, Mengen mit immer größerer Mächtigkeit erzeugen:

Sei M eine Menge. Wir definieren $\mathcal{P}^n(M)$ für $n \in \mathbb{N}$ rekursiv durch

$$\mathcal{P}^0(M) = M,$$

$$\mathcal{P}^{n+1}(M) = \mathcal{P}(\mathcal{P}^n(M)) \text{ für } n \in \mathbb{N}.$$

Dann gilt $|\mathcal{P}^n(M)| < |\mathcal{P}^{n+1}(M)|$ für alle $n \in \mathbb{N}$.

Sei weiter $M^* = \bigcup_{n \in \mathbb{N}} \mathcal{P}^n(M)$. Dann gilt $|\mathcal{P}^n(M)| < |\mathcal{P}^{n+1}(M)| \leq |M^*|$ für alle $n \in \mathbb{N}$. Durch die Vereinigung von

$M, \mathcal{P}(M), \mathcal{P}^2(M), \dots$

finden wir also eine Menge M^* von noch größerer Mächtigkeit. Wir können nun wieder $\mathcal{P}(M^*)$ bilden und haben $|M^*| < |\mathcal{P}(M^*)|$, usw. usw. Was hier genau „usw. usw.“ bedeutet, wird erst später klar werden, wenn wir die transfiniten Zahlen zur Verfügung haben.

Interpretation

Wir fassen die wichtigsten Ergebnisse noch einmal zusammen:

Es gibt Größenunterschiede im Reich des Unendlichen, wobei wir zwei beliebige Mengen als gleich groß ansehen, falls sie bijektiv aufeinander abbildbar sind. Weiter gilt, daß die Potenzmenge einer Menge immer von echt größerer Mächtigkeit ist als die Menge selbst.

Die natürliche Frage ist jetzt:

Um wieviel größer ist die Potenzmenge?

Diese Frage führt uns zum berühmtesten Problem der Mengenlehre, dem Cantorschen Kontinuumsproblem. David Hilbert, zu dieser Zeit der führende mathematische Kopf, hat 1900 auf dem internationalen Mathematikerkongreß in Paris die Frage *Um wieviel größer als $\mathbb{N}$ ist $\mathbb{R}$?* an die erste Stelle seiner Liste von 23 Jahrhundert-Problemen gestellt.

Wir formulieren das Problem im folgenden Kapitel genauer.

Der Beweis des Satzes von Cantor zeigt die Idee der Diagonalisierung in ihrer einfachsten und klarsten Form. Eine Funktion $f: M \rightarrow \mathcal{P}(M)$ enthält genügend Information, um ein $D \subseteq M$ zu konstruieren mit $D \notin \text{rng}(f)$. f kann also niemals surjektiv sein.

Wir schließen dieses Kapitel mit Cantors Originaldarstellung seines Verfahrens, zuerst vorgetragen bei der ersten Jahrestagung der DMV 1891.

Georg Cantors Diagonalverfahren

„In dem Aufsatz betitelt: Über eine Eigenschaft des Inbegriffs aller reellen algebraischen Zahlen ..., findet sich wohl zum ersten Male ein Beweis für den Satz, daß es unendliche Mannigfaltigkeiten gibt, die sich nicht gegenseitig eindeutig auf die Gesamtheit aller endlichen ganzen Zahlen $1, 2, 3, \dots, v, \dots$ beziehen lassen, oder, wie ich mich auszudrücken pflege, die nicht die Mächtigkeit der Zahlenreihe $1, 2, 3, \dots, v, \dots$ haben ...

Es läßt sich aber von jenem Satze ein viel einfacherer Beweis liefern ...

Sind nämlich m und w irgend zwei einander ausschließende Charaktere [z. B. 0 und 1], so betrachten wir den Inbegriff M von Elementen

$$E = (x_1, x_2, \dots, x_v, \dots),$$

welche von unendlich vielen Koordinaten $x_1, x_2, \dots, x_v, \dots$ abhängen, wo jede dieser Koordinaten entweder m oder w ist. M sei die Gesamtheit aller Elemente E [i.e. $M = {}^I\{m, w\}$, wobei $I = \mathbb{N} - \{0\}$].

... Ich behaupte nun, daß eine solche Mannigfaltigkeit M nicht die Mächtigkeit der Reihe $1, 2, \dots, v, \dots$ hat.

Dies geht aus folgendem Satze hervor:

„Ist $E_1, E_2, \dots, E_v, \dots$ irgendeine einfache unendliche Reihe von Elementen der Mannigfaltigkeit M , so gibt es stets ein Element E_0 von M , welches mit keinem E_v übereinstimmt.“

Zum Beweise sei

$$E_1 = (a_{1,1}, a_{1,2}, \dots, a_{1,v}, \dots),$$

$$E_2 = (a_{2,1}, a_{2,2}, \dots, a_{2,v}, \dots),$$

$$\dots \dots \dots$$

$$E_\mu = (a_{\mu,1}, a_{\mu,2}, \dots, a_{\mu,v}, \dots).$$

$$\dots \dots \dots$$

Hier sind die $a_{\mu,v}$ in bestimmter Weise m oder w . Es werde nun eine Reihe $b_1, b_2, \dots, b_v, \dots$, so definiert, daß b_v auch nur gleich m oder w und von $a_{v,v}$ *verschieden* sei.

Ist also $a_{v,v} = m$, so ist $b_v = w$, und ist $a_{v,v} = w$, so ist $b_v = m$.

Betrachten wir alsdann das Element

$$E_0 = (b_1, b_2, b_3, \dots)$$

von M , so sieht man ohne weiteres, daß die Gleichung

$$E_0 = E_\mu$$

für keinen positiven ganzzahligen Wert von μ erfüllt sein kann, da sonst für das betreffende μ und für alle ganzzahligen Werte von v

$$b_v = a_{\mu,v},$$

also auch im besonderen

$$b_\mu = a_{\mu,\mu}$$

wäre, was durch Definition von b_v ausgeschlossen ist. Aus diesem Satze folgt unmittelbar, daß die Gesamtheit aller Elemente von M sich nicht in die Reihenform: $E_1, E_2, \dots, E_v, \dots$ bringen läßt, da wir sonst vor dem Widerspruch stehen würden, daß ein Ding E_0 sowohl Element von M , wie auch nicht Element von M wäre.

Dieser Beweis erscheint nicht nur wegen seiner großen Einfachheit, sondern namentlich auch aus dem Grunde bemerkenswert, weil das darin befolgte Prinzip sich ohne weiteres auf den allgemeinen Satz ausdehnen läßt, daß die Mächtigkeiten wohldefinierter Mannigfaltigkeiten kein Maximum haben oder, was dasselbe ist, daß jeder gegebenen Mannigfaltigkeit L eine andere M an die Seite gestellt werden kann, welche von stärkerer Mächtigkeit ist als L ... “

(Georg Cantor 1892, „Über eine elementare Frage der Mannigfaltigkeitslehre“)

11. Die Kontinuumshypothese

$\mathbb{R}$ ist von größerer Mächtigkeit als $\mathbb{N}$. *Viel größer oder nur ein wenig größer?* – das ist die Frage. Die Cantorsche Kontinuumshypothese besagt, daß die Kluft zwischen $\mathbb{N}$ und $\mathbb{R}$ minimal ist, es also keine Mächtigkeiten gibt, die echt zwischen $|\mathbb{N}|$ und $|\mathbb{R}|$ liegen:

Kontinuumshypothese (CH)

Sei M eine Menge und es gelte $|\mathbb{N}| \leq |M| \leq |\mathbb{R}|$.

Dann gilt: $|\mathbb{N}| = |M|$ oder $|M| = |\mathbb{R}|$.

Anders formuliert:

Es gibt keine Menge M mit $|\mathbb{N}| < |M| < |\mathbb{R}|$.

Noch einmal anders formuliert:

Jede Teilmenge der reellen Zahlen ist abzählbar oder gleichmächtig zu den reellen Zahlen.

Übung

Zeigen Sie die Äquivalenz dieser drei Formulierungen.

Die Abkürzung (CH) steht für engl. „Continuum Hypothesis“ und ist mittlerweile allgemein gebräuchlich.

Besonders beim Betrachten der dritten Form fällt auf, wie einfach das Problem zu formulieren ist. Ähnlich wie manche klassische Probleme der Zahlentheorie läßt sich die Frage, die die Kontinuumshypothese stellt, auch einem interessierten Laien schnell erklären, im Gegensatz etwa zur Riemannschen Vermutung, die eine Aussage macht über die Nullstellen einer speziellen komplexwertigen Funktion ζ , der Riemannschen Zeta-Funktion (die Vermutung ist bis heute offen). Die Riemannsche Vermutung ist ebenso natürlich wie die Frage nach der Gültigkeit der Kontinuumshypothese, aber für den Nichtmathematiker schwerer zugänglich. Einfach zu formulierende Probleme sind zwar erfreulich, aber entscheidend ist letztendlich ihre Natürlichkeit innerhalb einer bereits als wertvoll erkannten mathematischen Umgebung, und diese Umgebung muß nicht immer leicht zu vermitteln sein. Manche Probleme drängen sich bei der Untersuchung eines mathematischen Gegenstandes regelrecht auf, sie entstehen aus intrinsischen Gründen, und die damit verbundene Dynamik trägt einen Großteil zur Entwicklung der Mathematik bei. Das Kontinuumsproblem ist ganz abgesehen von seiner einfachen Darstellbarkeit in dieser Hinsicht besonders natürlich: Die mathematische Umgebung bildet die Tatsache, daß es abzählbare und nicht-abzählbare Mengen gibt, und daß gerade *die* beiden Grundstrukturen $\mathbb{N}$ und $\mathbb{R}$

der Mathematik hier auseinanderfallen. Hat man dies einmal akzeptiert, so drängt sich die Frage nach der Größe der Kluft zwischen den natürlichen und den reellen Zahlen von selbst auf. Innerhalb der Ordinalzahltheorie der Mengenlehre kann man ein kanonisches Objekt ω_1 definieren, das minimal größer ist als $\mathbb{N}$: Es gilt $|\mathbb{N}| < |\omega_1|$, aber es existiert keine „Zwischenmenge“ M mit $|\mathbb{N}| < |M| < |\omega_1|$. Die Frage von (CH) lautet dann einfach: Gilt $|\mathbb{R}| = |\omega_1|$? Dieses Objekt ω_1 ist eine Grundstruktur der Mathematik, und in der mengentheoretischen Forschung steht es gleich neben $\mathbb{N}$ und $\mathbb{R}$. Wir definieren dieses Objekt im zweiten Abschnitt.

Nun denn: *Ist (CH) richtig oder falsch oder immer noch ungelöst?* Die Antwort ist beunruhigend. Sie wird gegeben durch den folgenden tiefen Satz:

Satz (*Fundamentalsatz der Mengenlehre*)

In der klassischen Mathematik gilt:

Die Kontinuumshypothese ist weder beweisbar noch widerlegbar.

Die klassische Mathematik ist hier ein Kunstausdruck (wie andernorts klassische Literatur oder Musik) und meint die durch die Tradition begründete und zur Zeit allgemein akzeptierte Mathematik. Genauer: Die durch eine übliche mengentheoretische Axiomatik des frühen 20. Jahrhunderts mengentheoretisch interpretierte Mathematik. Es ist die Mathematik, wie sie heute überall in Lehrbüchern zu finden ist. Es gibt zur Zeit einen sehr einheitlichen Bestand von allgemein anerkannten Methoden und Argumenten, und wir haben ihn hier überall verwendet, speziell etwa Induktion im Beweis des Satzes von Cantor-Bernstein und abstrakte Auswahl im Beweis des Vergleichbarkeitssatzes. Der Bestand wird, ebenso wie die Sprachstruktur innerhalb von Definition, Satz und Beweis, zumeist durch Vormachen und Nachahmung weitergegeben, und er wird oft nur in der mathematischen Logik explizit diskutiert.

Der Leser wird vielleicht sagen: *Es gibt doch nur eine Mathematik!* Damit dieser Satz Sinn hat, muß man sagen, was genau Mathematik ist. Und jede Definition, die dann *die* Mathematik liefert, ist ein Dogma, und kann von der mathematischen Praxis leicht überholt werden. Erfahrungsgemäß richtig ist: Mathematik als menschliche Tätigkeit ist eine ungemein streitfreie und zuverlässige Sache, und gerade die Streitfreiheit und Zuverlässigkeit meint man, wenn man die Mathematik vor anderen Wissenschaften herausheben möchte. Das soziale Phänomen ist es, das die Existenz *der einen* Mathematik suggeriert. Es herrscht Einigkeit darüber, ob ein vorgelegtes Ergebnis aus den und den Grundannahmen mit den und den logischen Schlüssen korrekt abgeleitet wurde, und Fehler in Argumenten werden von Kollegen normalerweise schnell entdeckt, und dann als solche auch mit „ich Esel“ und nicht mit „du Narr“ akzeptiert. (Daß dies etwas ganz Wunderbares ist, zeigt ein Vergleich mit der nicht unähnlich aufgebauten Juristerei.)

Normalerweise wird die Angabe der Grundannahmen und der logischen Schlußregeln unterdrückt, und dies geschieht einfach deswegen, weil sie in den meisten Fällen die gleichen sind. (Wir sagen auch nicht außerhalb der Mathematik: „Ich spreche jetzt deutsch in der üblichen Grammatik: Ich heiße Georg.“) Diese stillschweigenden Voraussetzungen bilden den klassischen Rahmen, und

die klassische Mathematik besteht aus den innerhalb dieses Rahmens erzielten oder generell erzielbaren Ergebnissen. Der Rahmen selber findet seinen mathematischen Ausdruck in einer formallogisch präsentierten axiomatischen Mengenlehre, die eines der sich sehr ähnlichen Axiomensysteme zugrunde legt, die in den ersten Jahrzehnten des 20. Jahrhunderts entwickelt worden sind. De facto genügt ein Bruchteil dieser axiomatischen Stärke für die meisten Disziplinen der Mathematik; der schon für die elementare Mengentheorie sicher nicht zu groß gewählte Rahmen erscheint der Analysis, Algebra, Geometrie, usw. fast schon als überdimensioniert.

Die Mengenlehre möchte aber dennoch ihrer Kardinalfrage auf den Grund gehen, und so unermeßlich weit das Objektmeer ihrer Basisaxiomatisierung der mathematischen Mitwelt auch erscheinen mag, liegen zu diesem Zweck doch Erweiterungen des Rahmens nahe. Neue Axiome braucht das Land, so rufen manche. Welche Axiome soll man nehmen? Müssen Axiome unmittelbar einleuchtend sein? Oder genügt es, wenn sie sich innerhalb einer sich entwickelnden Theorie herauskristallisieren und aufdrängen? Ist das Bild einer Verzweigung oberhalb der Grundaxiome das richtige? Ist diese Verzweigung uferlos, oder gibt es nur eine Handvoll sich widersprechender natürlicher und struktureicher Fortsetzungen?

Es gibt attraktive Erweiterungen des klassischen Rahmens, die (CH) positiv wie negativ entscheiden, und auch solche, die (CH) weiter offenlassen. Diese Axiome haben selbst unter Mengentheoretikern bislang keine allgemeine Akzeptanz gefunden. Inwieweit die Geschichte dafür verantwortlich ist, daß ein natürliches und starkes Axiom wie das Axiom der Konstruierbarkeit von Gödel nicht zum Kanon gehört, ist eine ebenso gewagte wie interessante Frage. Das Axiom besagt, daß die Ordinalzahlen den Kern der Mengenwelt bilden: Alles, was es gibt, windet sich in beschreibbarer Weise um dieses Rückgrat des Mengenuniversums herum. Dunkle Mengen existieren nicht, die Leuchtkraft der Ordinalzahlen erreicht jede Menge. Auswahlakte „ein ...“ können in wunderbarer Weise immer durch ein „das ...“ ersetzt werden, und es kommt noch besser: Die Kontinuumshypothese ist beweisbar, wenn man dieses Axiom, schamanistisch-postmodern „ $V = L$ “ genannt, akzeptiert, und sie ist beweisbar in einer Weise, daß einem fast die Krokodilstränen kommen. Man darf die These wagen, daß Cantor diesem Axiom bereitwillig die Tür geöffnet und es als den Gast begrüßt hätte, den er so schmerzlich vermißt hatte.

Was also ist schlecht am Axiom $V = L$? Nichts – aber es gibt Konkurrenz. Die Konkurrenz zu $V = L$ ist die Theorie der großen Kardinalzahlexiome, und bereits mittelstarke derartige Axiome brechen „ $V = L$ “ das Rückgrat: Sie beweisen, daß „ $V \neq L$ “ gilt. Weiter droht ständig das Subtheorieargument: Das, was unter „ $V = L$ “ alles war, erscheint nun als ein echter Teil von dem, was nun alles ist. Die Welt des – von Gödel selber abgelehnten – Gödelaxioms bleibt haargenau die gleiche, aber es gibt nun etwas außerhalb dieser Welt. Und es finden sich faszinierende Objekte und Strukturen in diesem Außerhalb.

Was also ist schlecht an großen Kardinalzahlexiomen? Nichts – aber auch diese Axiome sind unter Mengentheoretikern nicht allgemein als neuer Rahmen, als „wahr“, akzeptiert, wenn auch seit Jahrzehnten ein brennendes, von allen Seiten

geteiltes Interesse an diesen Axiomen besteht, und gute Argumente für die Erweiterung des Highways der Unendlichkeit vorliegen. Große Kardinalzahlaxiome entscheiden zwar (CH) nicht, aber sie entscheiden die Frage der Kardinalität von vielen Teilmengen von $\mathbb{R}$ zugunsten von (CH), und allgemeiner beweisen sie einen Satz von Axiomen, der als das Analogon der Dedekind-Peano Axiome der Zahlentheorie für die reellen Zahlen gelten darf. Weiter bilden sie einen babylonischen Turm, in dessen Stockwerken sich die Mengenlehre als Theorie komplett und in schöner Ordnung unterbringen läßt – sehr viel in der modernen Mengenlehre spielt sich außerhalb der logischen Kraft einer klassischen Axiomatik ab. „ $V = L$ “ kann übrigens mit einem Subtheorieargument kontern und sehr große Kardinalzahlen studieren, ohne sie jemals ganz zu besitzen. Das Argument erscheint nicht so natürlich wie das der Konkurrenz, aber es ist keineswegs völlig absurd; man kann es sogar als Synthese auffassen. Die Zukunft wird zeigen, ob eine Standardmengenlehre einmal „ $V = L$ “ oder große Kardinalzahlaxiome oder etwas ganz anderes enthalten wird, oder ob es bei der Verzweigung von interessanten Theorien oberhalb eines klassischen Kerns bleibt, wie die Situation zur Zeit wohl am neutralsten beschrieben wird.

Neben Erweiterungen sind auch ganz andere Szenarien denkbar: Eine Abschwächung des Rahmens aufgrund der Entdeckung eines Widerspruchs innerhalb der Rahmentheorie, und eine Umformulierung des Begriffs der reellen Zahlen und damit von (CH). Es könnte auch sein, daß jemand eine ganz andere Fundierung der Mathematik findet, die die Mengenlehre in einem neuen Licht erscheinen läßt. Wer weiß. Wir können nicht sagen, daß (CH) ein vages oder absolut unlösbares Problem ist. Wir können nur sagen: In der klassischen Mathematik gibt es keinen Beweis der Hypothese, und auch keinen Beweis ihrer Negation (es sei denn, die klassische Mathematik ist widersprüchlich, denn dann ist alles beweisbar). Wir werden dieses Resultat nun etwas genauer erläutern.

Unabhängigkeitsbeweise

Der Beweis des Fundamentalsatzes ruht auf zwei Säulen. Die eine stammt von Kurt Gödel (1938), die andere von Paul Cohen (1963). Gödel hat gezeigt: (CH) ist nicht widerlegbar, d. h. die Verneinung *non* (CH) der Kontinuumshypothese ist nicht beweisbar. Cohen (geb. 1934) hat gezeigt: (CH) ist nicht beweisbar.

Eine (in der klassischen Mathematik) weder beweisbare noch widerlegbare Aussage nennt man *unabhängig* (von der klassischen Mathematik). Daß es solche unabhängigen Aussagen gibt, wußte man bereits seit den Gödelischen Unvollständigkeitssätzen (1931). Allerdings werden diese Aussagen – mit einer Diagonalmethode! – sehr abstrakt konstruiert und ihr mathematischer Gehalt ist von einem ganz anderen Typ als (CH); sie besagen „ich bin nicht beweisbar“ oder „der zugrundeliegende Rahmen ist widerspruchsfrei“. Mit dem Beweis der Unabhängigkeit von (CH) hatte man zum ersten Mal eine unabhängige Aussage in Gestalt einer üblichen mathematischen Fragestellung gefunden.

Das Folgende hat beschreibenden Charakter und muß in vielen Punkten, die bei der Ausführung der hier dargestellten Ideen eine große Rolle spielen, ungenau und skizzenhaft bleiben. So nehmen wir etwa durchgehend die Widerspruchsfreiheit der üblichen Mathematik an, die sich aufgrund der Unvollständigkeitssätze von Gödel mathematisch nicht beweisen läßt. Die offizielle, aber hier zu umständliche Formulierung wäre etwa von der Form: Ist dieses und jenes axiomatische System widerspruchsfrei, so bleibt es widerspruchsfrei, wenn wir diese oder jene Aussage als neues Axiom zu dem System mit hinzunehmen. Weiter haben wir noch keinen formalen Rahmen entwickelt, der für die saubere Formulierung derartiger Resultate nötig wäre. Und auch dann hätte man noch einmal zwischen einer formalisierten Sprache, in der wir Mathematik betreiben, und ihrem kodierten Abbild innerhalb der mathematischen Objektwelt scharf zu trennen. Derlei Unterscheidungen sind für das Funktionieren der mathematischen Untersuchung der Mathematik selbst von großer Bedeutung, und die Verwechslung von Sprachebenen ist ein zeitloser Quell der Verwirrung. Für hier genügt uns eine Beschreibung, die dem Leser ungefähr ein Gefühl gibt, wie der Unabhängigkeitschase läuft. Um die ungemein subtilen Ausweichmanöver, die er zu vollführen hat, um nicht in die überall aufgestellten Fallen der logischen Unlauterkeit zu tappen, können wir uns hier nicht kümmern.

Wie sehen die beiden Säulen aus? Zunächst ist es nötig zu definieren, was „beweisbar im üblichen Rahmen“ heißt. Wie schon im zweiten Kapitel angedeutet, kann man einen formalen Beweisbegriff definieren im Sinne von „beweisbar mit Hilfe von bestimmten Axiomen (Grundannahmen) und einem genau definierten System aus Schlußregeln (Kalkül)“. Der „normale Beweise“ führende Mathematiker muß für seine Arbeit dieses formale System aus Kalkül und Axiomen gar nicht kennen, es ist eine Art Sekretär im Hintergrund, der jederzeit in der Lage ist, handschriftliche Notizen sauber und akkurat zu tippen; und diese Tätigkeit ist für „normale Beweise“ ziemlich überflüssig.

Für einen Unabhängigkeitsbeweis ist aber die Existenz eines solchen formalen Systems unerläßlich, denn hier ist die Rede davon, daß etwas nicht beweisbar oder widerlegbar ist, kein Beweis also eine bestimmte Frage beantwortet; und hierzu muß man offenbar festlegen, was man unter einem Beweis verstehen will.

Zum Glück müssen nun auch Mathematiker, die Unabhängigkeitsbeweise führen wollen, nicht das Dasein eines akkuraten Sekretärs fristen. Wichtig ist nur, daß nach der Aufstellung eines formalen Systems der Begriff „beweisbar“ ein mathematischer Begriff geworden ist, den man verwenden darf – wie er z. B. im Fundamentalsatz verwendet wird.

Von zentraler Bedeutung für einen Unabhängigkeitsbeweis ist nun der Begriff eines *Modells*, den wir hier kurz skizzieren wollen, und durch den der formal denkende und dienstbeflissene Sekretär wieder in den Hintergrund rückt – wo er auch hingehört.

Modelle

Ein Modell ist intuitiv eine Welt für ein mathematisches Axiomensystem, ein Bereich von Objekten, innerhalb dessen die Axiome gelten, oder etwas weniger hochgestochen, ein konkretes Beispiel. So sind etwa die Ebene und die Kugel zwei Modelle für gewisse Axiomensysteme der Geometrie. Ein Axiomensystem

ist dabei einfach eine Menge von mathematischen Aussagen in einer bestimmten formalen Sprache. Wir wollen uns wieder damit begnügen, daß eine solche formale Sprache im Hintergrund vorhanden ist, und formulieren Aussagen wie gehabt.

Wir betrachten etwa die Aussage

$\varphi =$ „Für alle x gibt es ein y mit $x + y = 0$.“

und die Modelle $\mathbb{N}$, $\mathbb{Z}$, $\mathbb{R}$. Wir interpretieren die Zeichen $+$, $=$, 0 für diese Modelle wie üblich. Bezogen auf $\mathbb{N}$ ist die Aussage offenbar falsch, bezogen auf $\mathbb{Z}$ und $\mathbb{R}$ ist sie richtig.

Ist nun $\mathcal{A} = \{\varphi_0, \varphi_1, \dots\}$ ein Axiomensystem, so heißt M ein Modell von $\mathcal{A}$, falls alle $\varphi \in \mathcal{A}$ bezogen auf M richtig sind. Ist φ bezogen auf ein Modell M richtig, so schreiben wir

$M \models \varphi$, gelesen „ M gilt φ “, „ M erfüllt φ “ oder „in M ist φ wahr“.

Wir schreiben $M \models \mathcal{A}$ für ein Modell M und ein Axiomensystem $\mathcal{A}$, falls $M \models \varphi$ für alle $\varphi \in \mathcal{A}$ gilt, und sagen „ M ist ein Modell von $\mathcal{A}$ “.

Trivial, aber wichtig ist:

(I) *In keinem Modell gilt zugleich φ und $\text{non } \varphi$.*

Dies ist die erste von zwei fundamentalen Eigenschaften eines Modellbegriffs. Die zweite besagt, daß der Modellbegriff logische Schlüsse respektiert:

Korrektheit des Modellbegriffs

Ist M ein Modell des Axiomensystems $\mathcal{A}$, und ist ψ eine Aussage, die sich mit Hilfe von $\mathcal{A}$ beweisen läßt (d. h. man darf für den Beweis die Aussagen von $\mathcal{A}$ als Hilfsmittel/Grundannahmen verwenden), so ist M auch ein Modell von ψ .

Kurz:

(II) *Gilt $M \models \mathcal{A}$ und ist ψ beweisbar mit Hilfe von $\mathcal{A}$, so gilt $M \models \psi$.*

Man nennt diese Aussage die *Korrektheit* der Gültigkeitsrelation oder des Modellbegriffs. Für die präzise Formulierung und den Beweis der Korrektheit ist wieder ein formaler Begriff von „Beweis“, „Aussage“, usw. unentbehrlich.

Bei der Definition des Modellbegriffs haben wir große Freiheit: Lediglich (I) und (II) stellen wir als Bedingungen; diese Bedingungen genügen für Unabhängigkeitsbeweise.

Es gibt einen kanonischen, auf Alfred Tarski (1933, 1935) zurückgehenden Modellbegriff, der genau der Intuition „ φ ist wahr bezogen auf M “ entspricht. Wir haben ihn im obigen Beispiel stillschweigend verwendet und werden ihn auch weiter stillschweigend verwenden. (Ein anderer Modellbegriff, den die Mengenlehre für Unabhängigkeitsbeweise verwendet, läßt z. B. neben „wahr“ oder „falsch“ die Elemente einer Booleschen Algebra als Wahrheitswerte zu.)

Modelle für die Mengenlehre

Man kann Axiomensysteme $\mathcal{A}$ für die Mengenlehre angeben, die unsere Intuition über den Mengenbegriff gut beschreiben, und uns erlauben, alle Konstruktionen, die wir bislang durchgeführt haben, zu rechtfertigen. (Wir werden in Abschnitt drei eine solche Axiomatisierung $\mathcal{A}$ angeben, die Zermelo-Fraenkel-Axiomatik ZFC, und alternative Systeme kurz diskutieren.) Viele Axiome postulieren die Existenz von bestimmten Objekten. So könnten etwa

„für alle x, y existiert die Menge $\{x, y\}$ “ und
 „für alle Mengen x existiert die Potenzmenge von x “

Elemente dieses Axiomensystems $\mathcal{A}$ sein. Sind sie nicht in $\mathcal{A}$ direkt enthalten, so werden sich diese Aussagen aber mit $\mathcal{A}$ beweisen lassen, wenn $\mathcal{A}$ ein umfassendes System für die Mengenlehre sein soll.

Es zeigt sich, daß in einem genügend reichhaltigen mengentheoretischen System $\mathcal{A}$ bereits die ganze klassische Mathematik interpretierbar ist, d. h.: Für jede mathematische Aussage Φ (der Algebra, der Analysis, der Mengenlehre, usw.) in der üblichen mathematischen Umgangssprache existiert eine Übersetzung ϕ von Φ in die formalisierte Sprache der Mengenlehre, für die gilt:

(+) Φ ist mathematisch beweisbar *gdw* ϕ ist beweisbar mit Hilfe von $\mathcal{A}$.

„Mathematisch beweisbar“ auf der linken Seite ist zu verstehen als „es existiert ein Beweis von Φ , wie er in Vorlesungen, Büchern, auf Tagungen, usw. geführt wird“. Zum Beispiel sind alle bisher geführten Beweise in diesem Text Beweise im Sinne der linken Seite. Auf der rechten Seite ist der akkurate Sekretär am Werk, der alles peinlichst genau aufschreibt.

„ Φ ist beweisbar (im Rahmen der üblichen Mathematik)“ heißt also nach (+): „ ϕ ist beweisbar mit Hilfe der Axiome $\mathcal{A}$ “, wobei $\mathcal{A}$ eine genügend reichhaltige Axiomatisierung der Mengenlehre darstellt.

Wir identifizieren im folgenden Φ und ϕ . Beweisbar in $\mathcal{A}$ ist dann nichts als „beweisbar in der klassischen Mathematik“ im Sinne der obigen Diskussion.

Zwischen Φ und ϕ gibt es noch eine, auch historisch vorhandene, Zwischenstufe, nämlich die der Interpretation einer mathematischen umgangssprachlichen Aussage Φ in eine mengensprachliche – noch nicht formalisierte – Aussage Φ' . Φ und Φ' verhalten sich etwa so wie ein Punkt P in einem zweidimensionalen Raum und ein $(x, y) \in \mathbb{R}^2$. Die Ebene kann in dieser Weise arithmetisch interpretiert werden. Griechische Dreiecke werden zu neuzeitlichen Teilmengen des $\mathbb{R}^2$. Analog kann man die Mathematik mengentheoretisch interpretieren, z. B. eine Funktion als eine Menge von geordneten Paaren auffassen, eine Gruppe als ein Paar $(G, \cdot)$ mit bestimmten Eigenschaften, usw. Nach einer gewissen Zeit identifiziert man die beiden Ebenen, und da die Interpretationsebene zumeist eine Spur genauer ist, möchte man sie bald nicht mehr missen, und vergißt sogar oft, daß die Übersetzung eine Übersetzung ist. Und so, wie man heute einen Punkt P der Ebene geradezu als Paar $(x, y) \in \mathbb{R}^2$ definiert wissen will, erwartet man von mathematischen Begriffen heute eine Formulierung innerhalb der Sprache der Mengenlehre. Vieles geht auch ohne solche Interpretationen: Die Griechen haben Geometrie auf hohem Niveau ohne eine arith-

metische Übersetzung betrieben, der Leser kannte Funktionen wahrscheinlich auch ganz gut, ohne Kuratowskipaare zu kennen, und allgemein sprach die Mathematik eine klare Sprache auch vor der Erfindung der Sprache der Mengenlehre. Die Funktionen der Analysis haben sich nicht geändert und die Sätze über Dreiecke sind die gleichen geblieben, sie werden aber heute anders behandelt und formuliert. Die höhere Genauigkeit ist der Grund, warum sich die mengentheoretische Interpretation durchgesetzt hat, auch wenn für viele Anwendungen diese Genauigkeit – und der Mengenreichtum – gar nicht gebraucht wird, ganz so, wie nicht für jede Berechnung die volle Prozessorleistung verwendet werden muß. Aber auch außerhalb der erhöhten Genauigkeit gibt es einen großen Gewinn: Eine universale Sprache für alle mathematischen Teilgebiete erlaubt erst einen uneingeschränkten Gedankenaustausch und eine gegenseitige Befruchtung der einzelnen Disziplinen. Und speziell für die Geometrie brachte die mengentheoretische Sprache eine Befreiung von der Arithmetik in Form der mengentheoretischen Topologie.

Die Beweisidee der Unabhängigkeit der Kontinuumshypothese

Der Fundamentalsatz wird nun der Grundidee nach wie folgt in zwei Teilen bewiesen (unter der Voraussetzung, daß $\mathcal{A}$ widerspruchsfrei ist):

- 1. Teil (Gödel)** Es gibt ein Modell M_1 von $\mathcal{A}$ mit $M_1 \models (\text{CH})$.
2. Teil (Cohen) Es gibt ein Modell M_2 von $\mathcal{A}$ mit $M_2 \models \text{non}(\text{CH})$.

Mit Hilfe von (+) und der Korrektheit des Modellbegriffs folgt dann, daß (CH) weder beweisbar noch widerlegbar ist (im Rahmen der üblichen Mathematik):

Wir nehmen an, (CH) sei mathematisch beweisbar, und betrachten das Modell M_2 von Cohen, in dem $\text{non}(\text{CH})$ gilt. Nach (+) ist (CH) beweisbar mit Hilfe von $\mathcal{A}$. Nach Korrektheit haben wir also $M_2 \models (\text{CH})$. Andererseits gilt $M_2 \models \text{non}(\text{CH})$ nach Wahl von M_2 . *Widerspruch*, denn nach (I) kann in einem Modell niemals eine Aussage und ihr Gegenteil zugleich wahr sein.

Übung

Argumentieren Sie analog mit Hilfe von M_1 , daß die Kontinuumshypothese (CH) nicht widerlegbar ist.

Ein Unabhängigkeitsbeweis besteht also darin, zwei Modelle zu konstruieren, in denen jeweils $\mathcal{A}$ gilt. In einem Modell soll zudem φ gelten, im anderen zudem $\text{non } \varphi$.

Für die erste Säule, den Beitrag von Gödel, startet man von einem großen Modell, und konstruiert innerhalb dieses Modells ein kleineres. Für die zweite Säule, den Beitrag von Cohen, erweitert man dagegen ein gegebenes Modell zu einem größeren Modell mit den gewünschten Eigenschaften. Die entsprechen-

den Beweistechniken heißen „innere Modelle“ und „forcing“ (Erzwingungsmethode). Die Methoden sind allgemein; mit ihrer Hilfe kann die Unabhängigkeit einer Vielzahl von Aussagen bewiesen werden, nicht nur die von (CH). Wir werden später ein weiteres Beispiel in der Suslin-Hypothese kennenlernen – benannt nach Mikhail Suslin (1894 – 1919).

Zum Schluß dieses Kapitels formulieren wir noch die natürliche Verallgemeinerung der Kontinuumshypothese.

Eine allgemeine Hypothese

Ersetzt man $\mathbb{R}$ in (CH) durch $\mathcal{P}(\mathbb{N})$, so hat die Kontinuumshypothese nur noch einen Parameter, nämlich $\mathbb{N}$, und in dieser Form ist dann die folgende Verallgemeinerung sehr naheliegend:

Verallgemeinerte Kontinuumshypothese (GCH)

Sei N eine unendliche Menge.

Sei M eine weitere Menge und es gelte $|N| \leq |M| \leq |\mathcal{P}(N)|$.

Dann gilt: $|N| = |M|$ oder $|M| = |\mathcal{P}(N)|$.

(GCH) steht für „Generalized Continuum Hypothesis“. Wir wissen vom Mächtigkeitssprung zwischen einer Menge und ihrer Potenzmenge, und (GCH) besagt, daß dieser Sprung so klein ist wie möglich für alle unendlichen Mengen. (GCH) ist offenbar falsch für endliche Mengen wegen $|\{0, 1\}| < |\{0, 1, 2\}| < |\mathcal{P}(\{0, 1\})|$.

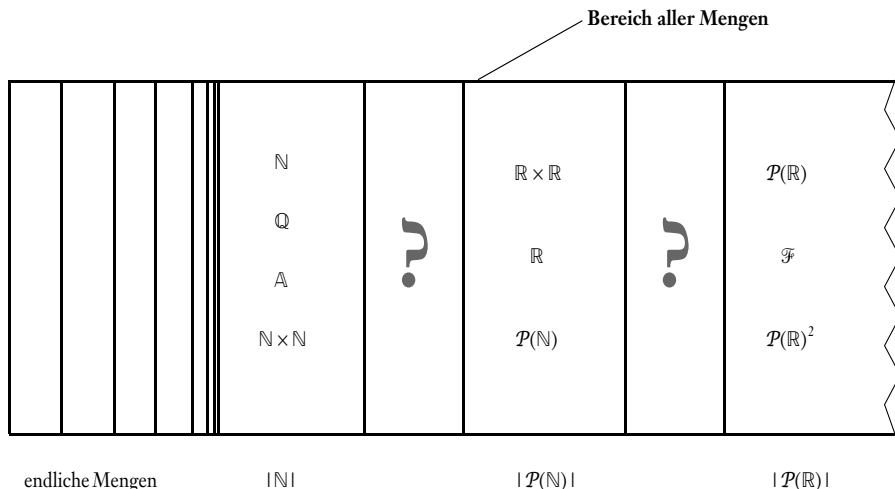
Im Fall $N = \mathbb{N}$ spezialisiert sich (GCH) wegen $|\mathcal{P}(\mathbb{N})| = |\mathbb{R}|$ zu (CH).

(GCH) ist ebenfalls unabhängig. Die Hypothese ist im Gödelschen Modell richtig, im Modell von Cohen ist sie trivialerweise falsch, da dort bereits (CH) nicht gültig ist.

Bei Cantor findet sich neben der Kontinuumshypothese (1878) lediglich die Vermutung, daß auch die Menge der reellen Funktion so klein ist wie möglich (1883b, §10), die Mengen $\mathbb{N}$, $\mathcal{P}(\mathbb{N})$, $\mathcal{P}(\mathcal{P}(\mathbb{N}))$ also die ersten drei unendlichen Mächtigkeiten repräsentieren. Ein allgemeines Interesse an höheren Mächtigkeiten setzte erst im ersten Viertel des 20. Jahrhunderts ein. Anwendungen, Formulierungen und Untersuchungen von (GCH) finden sich in [Hausdorff 1914, VI, §8] und [Lindenbaum / Tarski, 1926].

Versuch einer Visualisierung

Das folgende Diagramm gibt eine Zusammenfassung unserer Ergebnisse über die Größe von Mengen und eine anschauliche Fassung der Kontinuumshypothese. Wir denken uns den Bereich aller Mengen – das mengentheoretische Universum – in Abschnitte von Mengen gleicher Mächtigkeit eingeteilt, wobei die Mächtigkeiten von links nach rechts ansteigen.



Das Diagramm beginnt links mit den endlichen Mengen, die wir in N -viele Streifen „kein Element“, „genau ein Element“, „genau zwei Elemente“, usw. einteilen können. (Bereits die 1-Schicht ist uferlos groß, da für jedes Objekt x die Menge $\{x\}$ dieser Schicht angehört.) Der Rest der Mengenwelt besteht aus unendlichen Mengen. Unter ihnen bilden die abzählbaren Mengen, die Mengen der Mächtigkeit von N , die kleinste Schicht. Wir wissen, daß die Mengen der Mächtigkeit von R und weiter die Mengen der Mächtigkeit von $P(R)$ Schichten bilden, die weiter rechts liegen. Allgemein gelangen wir durch Anwendung der Potenzmengenoperation zu immer größeren Schichten. Die Aussage der Kontinuumshypothese und ihrer Verallgemeinerungen ist, daß diese Schichten aneinander lückenlos anschließen, daß also die mit einem Fragezeichen gekennzeichneten Bereiche leer sind.

Zur Beschreibung der Länge des Streifenbandes nach rechts brauchen wir die Ordinalzahlen. Denn auch nach allen Schichten

$$N, P(N), P^2(N) = P(P(N)), P^3(N), \dots$$

gibt es noch neue Schichten: Die Mengen der Mächtigkeit von $M = \bigcup_{n \in N} P^n(N)$ bilden eine Schicht hinter allen Schichten der Mengen $N, P(N), P^2(N)$, usw. (De facto ist M ein Repräsentant der auf alle $P^n(N)$ nächstfolgenden Schicht, wie wir später zeigen werden; das nächste Fragezeichen taucht also erst wieder beim Übergang von M zu $P(M)$ auf, und nicht etwa unmittelbar vor M .) Durch Bildung von

$$P(M), P^2(M), \dots, P^n(M), \dots, M' = \bigcup_{n \in N} P^n(M), P(M'), P^2(M'), \dots,$$

gelangen wir zu immer größeren Mächtigkeiten. Die Ordinalzahlen sind gerade die Kilometersteine dieses nie bis zu seinem Ende beschreitbaren, unbeschreiblich komplexen Weges nach rechts.

Georg Cantors erste Erwähnung der Kontinuumshypothese

„Da auf diese Weise für ein außerordentlich reiches und weites Gebiet von Mannigfaltigkeiten die Eigenschaft nachgewiesen ist, sich eindeutig und vollständig einer begrenzten, stetigen Geraden oder einem Teile derselben (unter einem Teile einer Linie jede in ihr enthaltene Mannigfaltigkeit von Punkten verstanden) zuordnen zu lassen, so entsteht die Frage, wie sich die verschiedenen Teile einer stetigen geraden Linie, d. h. die verschiedenen in ihr denkbaren unendlichen Mannigfaltigkeiten von Punkten hinsichtlich ihrer Mächtigkeit verhalten. Entkleiden wir dieses Problem seines geometrischen Gewandes und verstehen, wie dies bereits in §. 3 auseinandergesetzt ist, unter einer *linearen* Mannigfaltigkeit reeller Zahlen jeden denkbaren Inbegriff unendlich vieler, von einander verschiedener reeller Zahlen, so fragt es sich in *wie viel* und in welche Klassen die linearen Mannigfaltigkeiten zerfallen, wenn Mannigfaltigkeiten von gleicher Mächtigkeit in eine und dieselbe Klasse, Mannigfaltigkeiten von verschiedener Mächtigkeit in verschiedene Klassen gebracht werden. Durch ein Induktionsverfahren, auf dessen Darstellung wir hier nicht näher eingehen, wird der Satz nahe gebracht, daß die Anzahl der nach diesem Einteilungsprinzip sich ergebenden Klassen linearer Mannigfaltigkeiten eine endliche und zwar, daß sie gleich zwei ist.

Darnach würden die linearen Mannigfaltigkeiten aus zwei Klassen bestehen, von denen die erste alle Mannigfaltigkeiten in sich faßt, welche sich auf die Form: *functio ips. v* (wo *v* alle positiven Zahlen durchläuft) bringen lassen; während die zweite Klasse alle diejenigen Mannigfaltigkeiten in sich aufnimmt, welche auf die Form: *functio ips. x* (wo *x* alle reellen Werte ≥ 0 und ≤ 1 annehmen kann) zurückführbar sind. Entsprechend diesen beiden Klassen würden daher bei den unendlichen linearen Mannigfaltigkeiten nur zweierlei Mächtigkeiten vorkommen; die genauere Untersuchung dieser Frage verschieben wir auf eine spätere Gelegenheit.

Halle a. S., den 11. Juli 1877.“

(Georg Cantor 1878, „*Ein Beitrag zur Mannigfaltigkeitslehre*“)

12. Kardinalzahlen und ihre Arithmetik

In diesem Kapitel fassen wir zum ersten Mal für jede Menge M die Mächtigkeit oder Kardinalzahl $|M|$ von M als ein Objekt auf. Auf die Probleme einer genauen Definition haben wir in Kapitel 4 bereits hingewiesen. Die intuitive Begriffsbildung durch Abstraktion, wie sie Cantor vorgeschlagen hat, kann man nicht als extensionale mengentheoretische Definition auffassen. Der Leser wird aber wohl nach allem, was bisher geschah, mit dem Ausdruck „die Mächtigkeit der reellen Zahlen“ schon lange etwas anfangen können. Etwas locker kann man das Vorgehen dieses Kapitels so beschreiben: Schauen wir einfach mal, was passiert, wenn wir $|M|$ als Objekt zu lassen, und mit diesen Objekten rechnen wollen. Bei diesem Vorhaben verwenden wir Kardinalzahlen dann streng genommen nur als ein bequemes Notationssystem: $\alpha \cdot b = b \cdot \alpha$ schreibt sich viel einfacher und sieht viel besser aus als $|A \times B| = |B \times A|$.

Felix Hausdorff begnügte sich mit einem ähnlich „formalen Standpunkt“:

Hausdorff (1914): „Mengen eines Systems, die einer gegebenen Menge und damit auch untereinander äquivalent sind, haben etwas Gemeinsames, das im Falle endlicher Mengen die Anzahl der Elemente ist und das man auch im allgemeinen Falle die Anzahl oder Kardinalzahl oder Mächtigkeit nennt. Über die absolute Beschaffenheit dieses neu eingeführten Etwas kann man allerhand verschiedene Auffassungen hegen. G. Cantor definiert die Mächtigkeit einer Menge als den Allgemeinbegriff, der durch Abstraktion von der individuellen Beschaffenheit ihrer Elemente entsteht. B. Russell definiert sie geradezu als die Gesamtheit oder Klasse ‚aller‘ mit jener Menge äquivalenten Mengen; dies halten wir bei der uferlosen und antinomischen Beschaffenheit dieser Klasse für bedenklich. Wenn wir analoge Beispiele aus anderen Gebieten der Mathematik heranziehen, wird die gegenwärtige Situation nicht klarer; denn wenn wir kongruenten Punktpaaren eine gemeinsame ‚Entfernung‘, parallelen Geraden eine gemeinsame ‚Richtung‘, ähnlichen Figuren eine gemeinsame ‚Form‘ beilegen, so können ja diese Begriffe außerdem wirklich durch Strecken, Winkel oder Zahlen präzisiert werden. Andererseits könnte man den Begriff der Mächtigkeit freilich entbehren und alles auf die Betrachtung äquivalenter Mengen beschränken, worunter aber die Bequemlichkeit des Ausdrucks erheblich leiden würde. Übrigens ist darauf aufmerksam zu machen, daß die genannten Schwierigkeiten auch schon bei endlichen Mengen bestehen, wo es ja an verschiedenen Auffassungen des Zahlbegriffs durchaus nicht mangelt. Wir werden uns einfach auf den formalen Standpunkt stellen und sagen: einem System von Mengen A ordnen wir eindeutig ein System von Dingen α zu derart, daß äquivalenten Mengen und nur solchen dasselbe Ding entspricht, d. h. daß aus $A \sim B$ [$|A| = |B|$] immer $\alpha = b$ folgt und umgekehrt. Diese neuen Dinge oder Zeichen nennen wir Kardinalzahlen oder Mächtigkeiten; wir sagen: A hat die Mächtigkeit α , A ist von der Mächtigkeit α , α ist die Mächtigkeit von A , wohl auch (indem wir α als Zahlwort verwenden) A hat α Elemente.“

Diese Definition von Kardinalzahlen per Ritterschlag ließe sich mathematisch etwa so formulieren. Sei S eine Menge (intuitiv eine umfassende Menge, bei Hausdorff System genannt). Sei $K = S/\sim$ die Menge der Äquivalenzklassen der Relation „gleichmächtig“ auf S , d. h.

$$S/\sim = \{ \{ B \in S \mid A \text{ und } B \text{ sind gleichmächtig} \} \mid A \in S \}.$$

Wir definieren nun eine Funktion F auf $S/\sim$ durch:

$$F(X) = \text{„ein } A \in X\text{“} \quad \text{für } X \in S/\sim.$$

Für $A \in S$ können wir dann die Kardinalzahl oder Mächtigkeit $|A|$ von A definieren durch:

$$|A| = F(X_A), \text{ wobei } X_A \text{ das eindeutige } X \in S/\sim \text{ ist mit } A \in X.$$

Dieses Vorgehen liefert eine abstrakte, aber einwandfreie Definition des Kardinalzahlbegriffs für Elemente einer beliebig großen, aber fest gewählten Menge S . Für alle $A \in S$ ist $|A|$ definiert, und $|A|$ ist zudem ein Objekt, das mit A gleichmächtig ist. (Verzichten wir hierauf, kann man auch $|A| = X_A$ setzen.)

Da S beliebig groß gewählt und gegebenenfalls erweitert werden kann, scheint dies ein brauchbares Vorgehen für alle praktischen Belange zu sein. Es ist aber alles andere als schön, eher eine Notlösung: Welche Mengen den Ritterschlag *Kardinalzahl* bekommen, bleibt unklar, da Kardinalzahlen über ein abstraktes „ein ...“ erhalten worden sind. Und $|A|$ ist immer nur für bestimmte Mengen definiert, nämlich für $A \in S$. („ $S = \text{Alles}$ “ ist nicht möglich, wie wir in Abschnitt 13 sehen werden.) Zudem ist das Vorgehen auch praktisch nicht unproblematisch, da wir mit Kardinalzahlen frei operieren wollen, und es geht ja nicht zuletzt um die Freiheit und „Bequemlichkeit des Ausdrucks“. Soll bei jeder Operation wieder eine Kardinalzahl herauskommen, so müssen wir prüfen, ob S reichhaltig genug war, ein Ergebnis zur Verfügung zu stellen. Für kleine Operationen wie Summe und Produkt zweier Kardinalzahlen wäre dies leicht zu sichern, aber ein Produkt über sehr viele $|A|$ werden wir i. a. nicht ausführen können, obwohl sich eine natürliche Definition für Produkte mit beliebig vielen Faktoren anbieten wird. Außerdem wird der nimmermüde Geist fragen: Was ist $|S|$?

Diese mengentheoretische Interpretation des Hausdorffschen Zeichensystems liefert also wieder keine befriedigende Definition. Wir müssen das obige Zitat, wohl weitestgehend der Intention Hausdorffs entsprechend, den Beschreibungen von Cantor und Fraenkel aus Kapitel 4 an die Seite (oder gegenüber) stellen. Diese Beschreibungen fördern die Intuition und den Blick fürs Wesentliche unheimlich, lösen aber das definitorische Problem nicht.

Es bleibt der Verweis auf die prinzipielle Entbehrlichkeit von Kardinalzahlen. Die prinzipiell mögliche Rückübersetzung von Aussagen mit Kardinalzahlen in relationale Aussagen, die nur $|A| = |B|$, $|A| \leq |B|$, usw. enthalten, ist es, die uns rettet, und die Aufnahme einer definitorischen Hypothek akzeptabel macht.

Die Ergebnisse des Folgenden werden dann also auch bei allen Vorbehalten gegen $|M|$ als Objekt, die der streng klösterliche Blick geltend machen kann, selbst vom Papst gebilligt werden müssen; an der mathematischen Keuschheit der Inhalte besteht kein Zweifel, auch wenn ihre Darstellung etwas sinnlichere Züge annimmt.

Es ist interessant, daß Hausdorff seinen „formalen Standpunkt“ in der zweiten, vielfach umgeschriebenen Auflage seines Buches sogar noch verschärft hat. Insbesondere fällt die Einschränkung auf ein System ganz weg (wie auch schon in der 1914-Fassung an späterer Stelle bei einer analogen Definition des „Ordnungstypus“ einer Menge), und es gibt einen neuen, elegant formulierten Zusatz eines Autors, der sein eigenes Werk überarbeitet, und auf ein notorisches Problem etwas genervt reagiert:

Hausdorff (1927): „... D. h. wir ordnen jeder Menge A ein Ding α zu derart, daß äquivalenten Mengen und nur solchen dasselbe Ding entspricht:

$$\alpha = \beta \quad \text{soviel wie} \quad A \sim B.$$

Diese neuen Dinge nennen wir Kardinalzahlen oder Mächtigkeiten...

Diese formale Erklärung sagt, was die Kardinalzahlen sollen, nicht was sie sind. Prägnantere Bestimmungen sind versucht worden, aber sie befriedigen nicht und sind auch entbehrlich. Relationen zwischen Kardinalzahlen sind nur ein bequemer Ausdruck für Relationen zwischen Mengen: das ‚Wesen‘ der Kardinalzahl zu ergründen müssen wir der Philosophie überlassen.“

Es gab 1927 bereits eine befriedigende Definition innerhalb der Mengenlehre, ganz ohne Philosophie. Wir werden sie im zweiten Abschnitt diskutieren.

Kardinalzahlen

Hier nun also, nach hoffentlich ausreichendem Vorspiel auf dem Theater, die definitorische Hypothek dieses ersten Akts:

Definition (*Kardinalzahlen*)

Die Mächtigkeiten von Mengen heißen *Kardinalzahlen*.

Die Kardinalzahl einer Menge M wird mit $|M|$ bezeichnet.

$|M|$ heißt die *Kardinalität* oder *Mächtigkeit* von M.

„Mächtigkeit“ verwendet Cantor relational für zwei Mengen seit etwa 1877, die in Kapitel 4 wiedergegebene Definition steht ganz zu Beginn seiner 1878 erschienenen Arbeit. Einen selbständigen Mächtigkeitsbegriff, der die natürlichen Zahlen umfaßt, diskutiert er ausführlich 1882. Am Ende der Arbeit von 1878 deutet sich ein selbständiger Begriff der „Mächtigkeit einer Menge“ bereits an. Entnommen hat Cantor das Wort einer Vorlesung von Jacob Steiner (1796 – 1863):

Cantor (1882b): „Den Ausdruck ‚Mächtigkeit‘ habe ich J. Steiner entlehnt *) [*] = M. s. dessen Vorlesung über synthetische Geometrie der Kegelschnitte, herausgeg. von Schröter; § 2.] der ihn in einem ganz speziellen, immerhin jedoch eng verwandten Sinne gebraucht, um auszusprechen, daß zwei Gebilde durch *projektivische* Zuordnung [bijektiv aufeinander abgebildet werden können].“

Gleichwertig zu „Mächtigkeit einer Menge“ verwendet Cantor ab 1887 „Kardinalzahl einer Menge“, definiert als „universale“ (vgl. 1.4). Der Gedanke an ranghohe Geistliche liegt nahe. Die Kardinäle selber haben ihren Namen von lateinisch *cardo*, Türangel, was

sich als metaphorischer Kaffeesatz für „an einer wichtigen Stelle“ anbietet, und als *cardinalis* dann ein Wort für *grundlegend, wichtig* liefert. (Siehe Duden 7 für Details.) Kardinalzahlen sind also wahrhaft Dreh- und Angelpunkte der Mengenlehre!

Traditionell werden kleine Fraktur-Buchstaben $\alpha, \beta, \gamma, \dots$ für Kardinalzahlen verwendet. (Der mit Stift und Papier bewehrte Leser kann sich ohne Verlust mit $a, b, c, \dots$ begnügen.)

Für einige spezielle Kardinalzahlen gibt es spezielle natürliche und daneben auch traditionelle Bezeichnungen. Zunächst zeichnen wir die natürlichen Zahlen als Kardinalitäten aus:

Definition (die Kardinalitäten $n \in \mathbb{N}$)

Für alle $n \in \mathbb{N}$ wird die Kardinalität von $|n|$ mit n bezeichnet.

Jedes $n \in \mathbb{N}$ heißt eine *endliche Kardinalität*.

Die anderen Kardinalitäten heißen *unendliche Kardinalitäten*.

Die Definition bevorzugt $\mathbb{N}$ -Endlichkeit gegenüber der Dedekind-Endlichkeit, was im Umfeld von Zahlen häufig geschieht.

Jedes $n \in \mathbb{N}$ ist damit eine Kardinalzahl. Es gilt $|M| = n$ genau dann, wenn $|M| = |n|$ gilt, d. h. genau dann, wenn M genau n Elemente besitzt.

Für die abzählbaren Mengen hat Cantor den ersten Buchstaben des hebräischen Alphabets, verziert mit dem Index Null, gewählt:

Definition (die Kardinalität $\aleph_0$)

Die Kardinalität von $\mathbb{N}$ wird mit $\aleph_0$ [aleph 0] bezeichnet.

Cantor (1895): „§6. Die kleinste transfinite Kardinalzahl Alef-null.

Die Mengen mit endlicher Kardinalzahl heißen ‚*endliche Mengen*‘, alle anderen wollen wir ‚*transfinite Mengen*‘ und die ihnen zukommenden Kardinalzahlen ‚*transfinite Kardinalzahlen*‘ nennen.

Die Gesamtheit aller *endlichen Kardinalzahlen* v bietet uns das nächstliegende Beispiel einer transfiniten Menge; wir nennen die ihr zukommende Kardinalzahl (§1) ‚*Alef-null*‘, in Zeichen $\aleph_0$, definieren also

$$\aleph_0 = \{v\}. \quad [\aleph_0 = |\{v \mid v \in \mathbb{N}\}|]$$

Die Zahl $\aleph_0$ ist größer als jede endliche Zahl...

Andererseits ist $\aleph_0$ die kleinste transfinite Kardinalzahl...

Diese von Cantor geliebte Bezeichnung hat eine gewisse mystische Kraft und ist weithin bekannt geworden. Der Leser wird vielleicht die Erzählung „Das Aleph“ des argentinischen Schriftstellers Jorge Luis Borges kennen. (Wer Borges nicht kennt, kennt ihn doch: Er steht hinter Jorge von Burgos, dem Bibliothekar in Umberto Ecos „Name der Rose“.) In „Das Aleph“ heißt es etwa: „Ich stieg heimlich hinunter, rutschte auf der verbotenen Treppe aus, fiel hin. Als ich die Augen aufschlug, sah ich das Aleph.“ Und zum Zeichen $\aleph$ selbst: „...auch wurde gesagt, daß das Aleph die Gestalt eines Menschen habe, der auf den Himmel und die Erde zeigt, um anzudeuten, daß die untere Welt Spiegel und Kartenbild der oberen sei.“ Mathematisch ist das Objekt $\aleph_0$ die heißumkämpfte und lange verbotene Schnittstelle zwischen Endlichkeit und Unendlichkeit. Sobald man es sieht, gibt es kein Halten mehr, und sobald

man es hat, beginnt die Mengenlehre. In dieser ist dann auch das Spiegel-Bild oft formuliert worden: Wir übertragen unsere endliche Intuition über Mengen so weit wie möglich auf den unendlichen Bereich.

Höher indizierte Alephs und einige weitere hebräische Buchstaben werden später noch Verwendung finden. Bislang unbenutzt in der Mengenlehre blieben die Zeichen des Sanskrit ...

Es gilt $|M| = \aleph_0$ für jede abzählbar unendliche Menge.

Eine traditionelle Bezeichnung für $|\mathbb{R}|$ ist $\mathfrak{c}$, was an „Continuum“ erinnert. Wir folgen dieser Tradition weitgehend, verwenden aber $\mathfrak{c}$ auch, insbesondere im Umfeld von $\aleph$ und $\mathfrak{b}$, als Zeichen für eine beliebige Kardinalität.

Definition (*< und $\leq$ für Kardinalzahlen*)

Seien $\alpha, \mathfrak{b}$ Kardinalzahlen, und seien A, B Mengen mit $|A| = \alpha, |B| = \mathfrak{b}$.

Wir definieren:

$$\alpha \leq \mathfrak{b} \text{ falls } |A| \leq |B|,$$

$$\alpha < \mathfrak{b} \text{ falls } |A| < |B|.$$

Arithmetische Operationen mit Kardinalzahlen

Für Kardinalzahlen fließen neben dem kleiner und kleinergleich auch die Definitionen der Addition, Multiplikation und Exponentiation ohne Mühe aus der Feder. Dies geschieht wie erwartet derart, daß die üblichen arithmetischen Operationen auf den natürlichen Zahlen, die ja nun per definitorischem Dekret zu Kardinalitäten ernannt worden sind, mit den neudefinierten Operationen, eingeschränkt auf die endlichen Kardinalzahlen, zusammenfallen.

Definition (*Addition, Multiplikation und Exponentiation von Kardinalzahlen*)

Seien $\alpha, \mathfrak{b}$ Kardinalzahlen, und seien A, B Mengen mit $|A| = \alpha, |B| = \mathfrak{b}$.

Wir definieren die *Summe* $\alpha + \mathfrak{b}$, das *Produkt* $\alpha \cdot \mathfrak{b}$ und die *Exponentiation* $\alpha^{\mathfrak{b}}$ von α und $\mathfrak{b}$ wie folgt:

$$\alpha + \mathfrak{b} = |A \times \{0\} \cup B \times \{1\}|,$$

$$\alpha \cdot \mathfrak{b} = |A \times B|,$$

$$\alpha^{\mathfrak{b}} = |{}^B A|.$$

Die Konstruktion $A \times \{0\} \cup B \times \{1\}$ in der Addition hat folgenden Grund: Es gilt $|A| = |A \times \{0\}|$, $|B| = |B \times \{1\}|$, und $A \times \{0\} \cap B \times \{1\} = \emptyset$. Wir verwenden also disjunkte Mengen der entsprechenden Kardinalitäten zur Addition. Eine Definition von $\alpha + \alpha = |A \cup A| = |A|$ ist sicher nicht das, was wir wollen. Für die wie oben definierte Addition gilt aber gewiß: Sind A, B disjunkte Mengen, so ist $|A \cup B| = |A| + |B|$.

Es ist klar nach all dem, was wir in den vorangehenden Kapiteln gezeigt haben, daß die Operationen wohldefiniert sind, d. h. verwenden wir andere Mengen

M, N mit $|M| = |A| = \alpha$ und $|N| = |B| = b$, so erhalten wir dieselben Ergebnisse für $\alpha + b$, $\alpha \cdot b$ und α^b .

Man zeigt leicht, daß die Operationen mit endlichen Kardinalitäten $n, m \in \mathbb{N}$ die üblichen Operationen auf den natürlichen Zahlen liefern. Der Leser überlege sich statt einer langweiligen und länglichen Übung vielleicht die folgenden Spezialfälle: $0 \cdot n = 0$, $n^0 = 1$ für alle $n \in \mathbb{N}$, $0^n = 0$ für alle $n \in \mathbb{N} - \{0\}$. [Für die leere Menge gilt $\emptyset: \emptyset \rightarrow M$ für alle Mengen M . Daher die 1 in n^0 . Allgemein gültige Sonderfälle sind $\alpha^0 = 1$, $1^\alpha = 1$ für alle α und $0^\alpha = 0$ für alle $\alpha \neq 0$.]

Für die Addition und die Multiplikation von beliebigen Kardinalzahlen gelten die aus dem Reich des Endlichen vertrauten Rechenregeln:

Übung

Die Operationen $+$, $\cdot$ sind kommutativ, assoziativ, und es gelten die Distributivgesetze, d. h. für alle Kardinalzahlen α, b, c gilt:

- (i) $\alpha + b = b + \alpha$, $\alpha \cdot b = b \cdot \alpha$,
- (ii) $(\alpha + b) + c = \alpha + (b + c)$, $(\alpha \cdot b) \cdot c = \alpha \cdot (b \cdot c)$,
- (iii) $(\alpha + b) \cdot c = \alpha \cdot c + b \cdot c$, $\alpha \cdot (b + c) = \alpha \cdot b + \alpha \cdot c$.

Übung

Seien α, b, c Kardinalzahlen, und es gelte $b \leq c$. Dann gilt:

- (i) $\alpha + b \leq \alpha + c$,
- (ii) $\alpha \cdot b \leq \alpha \cdot c$,
- (iii) $\alpha^b \leq \alpha^c$, falls $b \neq 0$ oder $\alpha \neq 0$.
- (iv) $b^a \leq c^a$.

Dagegen bleibt ein $<$ i. a. nicht erhalten: $0 < 1$, aber $0 + \aleph_0 = \aleph_0 = 1 + \aleph_0$.

Altes in neuem Gewande

Wir stellen in der neuen Notation einige Resultate zusammen, die wir in den vorangehenden Kapiteln (incl. der Übungen) bewiesen haben.

Für $\aleph_0 = |\mathbb{N}|$, $c = |\mathbb{R}|$, $\mathfrak{f} = |\mathfrak{S}|$ gilt:

- (i) $\aleph_0 < c < \mathfrak{f}$,
- (ii) $\aleph_0 + \aleph_0 = \aleph_0 \cdot \aleph_0 = \aleph_0$,
- (iii) $c = c + c = c \cdot c$,
- (iv) $\mathfrak{f} = \mathfrak{f} + \mathfrak{f} = \mathfrak{f} \cdot \mathfrak{f}$,
- (v) $c = 2^{\aleph_0} = \aleph_0^{\aleph_0} = c^{\aleph_0}$,
- (vi) $\mathfrak{f} = 2^c = \aleph_0^c = c^c = \mathfrak{f}^c$.

Für alle Mengen M gilt: $|\mathcal{P}(M)| = 2^\alpha$ falls $|M| = \alpha$.

Unsere Hauptsätze schreiben sich nun sehr elegant, denn für alle Kardinalzahlen α, b haben wir:

- (i) $\alpha \leq \beta$ und $\beta \leq \alpha$ folgt $\alpha = \beta$, (Satz von Cantor-Bernstein)
- (ii) $\alpha \leq \beta$ oder $\beta \leq \alpha$, (Vergleichbarkeitssatz)
- (iii) $\alpha < 2^\alpha$. (Satz von Cantor)

Für alle unendlichen Kardinalzahlen α gilt weiter:

- (i) $\aleph_0 \leq \alpha$,
- (ii) $\alpha + 1 = \alpha$,
- (iii) $\alpha + \aleph_0 = \alpha$,
- (iv) $\alpha + \alpha = \alpha$ folgt $2^\alpha \cdot 2^\alpha = 2^\alpha$.

Die neuen Kardinalzahlen führen nun nicht nur zu kompakten Notationen, sondern suggerieren Rechengesetze, mit denen sich einige Resultate innerhalb einer Zeile beweisen lassen, etwa $|\mathbb{R}^2| = |\mathbb{R}|$. Dies verdient einen eigenen Zwischenabschnitt.

Rechenregeln der Exponentiation

Dem Leser wird die folgende Übung dringend ans Herz gelegt:

Übung

Seien α, β, γ Kardinalzahlen. Dann gilt:

- (i) $\alpha^{\beta+\gamma} = \alpha^\beta \cdot \alpha^\gamma$,
- (ii) $(\alpha \cdot \beta)^\gamma = \alpha^\gamma \cdot \beta^\gamma$,
- (iii) $(\alpha^\beta)^\gamma = \alpha^{\beta \cdot \gamma}$.

Mit diesen einfachen Rechenregeln gewinnen wir viele Mächtigkeitsresultate von früher sehr einfach. Einige Beispiele sind:

- (i) $|\mathbb{R}^2| = |\mathbb{R} \times \mathbb{R}| = 2^{\aleph_0} \cdot 2^{\aleph_0} = 2^{\aleph_0 + \aleph_0} = 2^{\aleph_0} = |\mathbb{R}|$,
- (ii) $|\mathbb{N}^{\mathbb{R}}| = (2^{\aleph_0})^{\aleph_0} = 2^{\aleph_0 \cdot \aleph_0} = 2^{\aleph_0} = |\mathbb{R}|$,
- (iii) $|\mathbb{R}^{\mathbb{R}}| = (2^{\aleph_0})^{2^{\aleph_0}} = 2^{\aleph_0 \cdot 2^{\aleph_0}} = 2^{2^{\aleph_0}} = 2^{|\mathbb{R}|} = |\mathcal{P}(\mathbb{R})|$.

Wir verwenden hier neben $|\mathbb{R}| = |\mathcal{P}(\mathbb{N})| = 2^{\aleph_0}$ lediglich die Gleichungen $\aleph_0 + \aleph_0 = \aleph_0$ und $\aleph_0 \cdot \aleph_0 = \aleph_0$. Hieraus folgt dann weiter:

$$2^{\aleph_0} \leq \aleph_0 \cdot 2^{\aleph_0} \leq 2^{\aleph_0} \cdot 2^{\aleph_0} = 2^{\aleph_0 + \aleph_0} = 2^{\aleph_0},$$

also die in (iii) verwendete Gleichung $\aleph_0 \cdot 2^{\aleph_0} = 2^{\aleph_0}$.

Die Beweise der Sätze über die Mächtigkeit mehrdimensionaler Kontinua haben durch diese Algebraisierung zwar erheblich an Kürze und Eleganz gewonnen, gegenüber den Beweisen des ersten Abschnitts dafür aber Kreativität und Anschauung eingebüßt. Für die abstrakte Kardinalzahlarithmetik, einem der Hauptanliegen auch der heutigen Mengenlehre, ist dieser simple „ 2^α -Kalkül“ aber unentbehrlich.

Allgemeine Summen und Produkte

Der Leser kann den „technischen Rest“ dieses Kapitels gefahrlos überspringen. Höhepunkte des Folgenden sind vielleicht die Ungleichung von König-Zermelo und der Multiplikationssatz. Wir diskutieren zudem einige Fragen, die man zur Arithmetik mit Kardinalzahlen stellen kann, und motivieren damit die allgemeine Theorie geordneter Mengen des zweiten Abschnitts. Das Folgende erscheint dort noch einmal in einem sehr milden Licht, das der Anschauung und den Photographien des Gedächtnisses sehr entgegenkommt.

Die richtigen Begriffe für bestimmte Probleme zu entwickeln und zu verwenden hat Vorrang. Daneben steht aber ein gewisses puristisches Interesse, die Kardinalzahlarithmetik zunächst ohne den Wohlordnungsbegriff zu entwickeln. Wie und daß dies geht, erscheint nicht uninteressant. Der Leser hat dann zudem die Möglichkeit, die beiden Methoden einander gegenüberzustellen.

Kardinalzahlen lassen sich nicht nur in Paaren oder endlich oft, sondern in beliebiger Menge sinnvoll addieren und multiplizieren. Für die Multiplikation brauchen wir zunächst einen allgemeinen Produktbegriff.

Definition (*allgemeines Kreuzprodukt*)

Sei I eine Menge, und seien A_i Mengen für $i \in I$.

Dann ist $\times_{i \in I} A_i$, das allgemeine Produkt der Mengen A_i , $i \in I$, definiert durch:

$$\times_{i \in I} A_i = \{ f \mid f: I \rightarrow \bigcup_{i \in I} A_i, f(i) \in A_i \text{ für alle } i \in I \}.$$

Das Produkt ist also die Menge der Transversalfunktionen f , die auf der Indexmenge I definiert sind und an jeder Stelle $i \in I$ einen Wert in A_i annehmen. Wir können $A_0 \times A_1$ zumeist gefahrlos mit $\times_{i \in \{0,1\}} A_i$ identifizieren, obwohl streng genommen $A_0 \times A_1$ eine Menge von geordneten Paaren (a_0, a_1) ist, $\times_{i \in \{0,1\}} A_i$ dagegen eine Menge von Funktionen der Form $\{ (0, a_0), (1, a_1) \}$.

Will man als „warm-up“ zeigen, daß das Produkt unendlich vieler nichtleerer Mengen immer nichtleer ist, so wird man bei der Konstruktion einer Transversalfunktion feststellen, daß eine (triviale) Definition der Form „ein ...“ gebraucht wird. (Generell gilt, daß in der Kardinalzahlarithmetik abstrakte Auswahlgeneratoren fast überall am Werk sind.) Es stellt sich dann – nach diesem ersten Resultat über die Existenz *einer* Transversalfunktion – der Intuition entsprechend heraus, daß es für unendliche Systeme aus A_i 's mit mehr als einem Element massenweise Transversalfunktionen gibt. Vgl. hierzu die Übung unten.

Definition (*Summe und Produkt von Kardinalzahlen*)

Sei I eine Menge, und seien α_i Kardinalzahlen für $i \in I$.

Weiter seien A_i Mengen für $i \in I$ mit $|A_i| = \alpha_i$.

Wir definieren die *Summe* $\sum_{i \in I} \alpha_i$ und das *Produkt* $\prod_{i \in I} \alpha_i$ der Kardinalzahlen α_i , $i \in I$, wie folgt:

$$\sum_{i \in I} \alpha_i = \left| \bigcup_{i \in I} A_i \times \{i\} \right|,$$

$$\prod_{i \in I} \alpha_i = \left| \times_{i \in I} A_i \right|.$$

Die Summe ist also allgemein definiert über die Kardinalität einer disjunkten Vereinigung, und das Produkt über die Menge der Transversalfunktionen, die durch die indizierten Mengen laufen.

Es ist klar, daß diese allgemeinen Summen- und Produktdefinitionen für endlich viele Summanden und Faktoren mit den alten übereinstimmen. So gilt zum Beispiel $\alpha_0 \cdot \alpha_1 \cdot \alpha_2 = \prod_{i \in \{0, 1, 2\}} \alpha_i$.

Übung

Seien I, J Mengen, und seien α_i, b_j Kardinalzahlen für $i \in I$ bzw. $j \in J$.

Dann gilt:

- (i) $\prod_{i \in I} \alpha_i \neq 0$ gdw $\alpha_i \neq 0$ für alle $i \in I$,
- (ii) $\sum_{i \in I} \alpha_i \neq 0$ gdw $\alpha_i \neq 0$ für ein $i \in I$,
- (iii) $\sum_{i \in I} \alpha_i \leq \prod_{i \in I} \alpha_i$ falls $\alpha_i \geq 2$ für alle $i \in I$,
- (iv) $\sum_{i \in I} \alpha_i \leq \sum_{j \in J} b_j$ falls ein injektives $f: I \rightarrow J$ existiert mit $\alpha_i \leq b_{f(i)}$ für alle $i \in I$,
- (v) $\prod_{i \in I} \alpha_i \leq \prod_{j \in J} b_j$, falls $b_j \neq 0$ für alle $j \in J$, und ein injektives $f: I \rightarrow J$ existiert mit $\alpha_i \leq b_{f(i)}$ für alle $i \in I$.

Es folgt, daß $\sum_{i \in I} \alpha_i = \sum_{i \in I} \alpha_{f(i)}$ für alle Bijektionen $f: I \rightarrow I$ gilt, und ein ebenso allgemeines Kommutativgesetz gilt für die Multiplikation. Nach Definition haben wir $\prod_{i \in I} \alpha_i = 1$ für $I = \emptyset$.

Es gelten weiter andere aus der endlichen Arithmetik bekannte Rechenregeln:

Übung

Sei I eine Menge, und seien α_i Kardinalzahlen für $i \in I$.

- (i) $\prod_{i \in I} \alpha_i^b = (\prod_{i \in I} \alpha_i)^b$ für alle Kardinalzahlen b ,
- (ii) $\prod_{i \in I} \alpha_i = \prod_{X \in Z} \prod_{i \in X} \alpha_i$ und
 $\sum_{i \in I} \alpha_i = \sum_{X \in Z} \sum_{i \in X} \alpha_i$

für jede Zerlegung Z von I in paarweise disjunkte Mengen.

Für die Unersättlichen seien schließlich auch noch die Distributivgesetze in ihrer allgemeinsten Form notiert:

Übung (allgemeine Distributivgesetze für Mengen und Kardinalzahlen)

Sei K eine Menge, und seien I_k Mengen für $k \in K$.

Weiter seien $A_{k,i}$ Mengen für $i \in I_k, k \in K$, und es sei

$\alpha_{k,i} = |A_{k,i}|$ für $i \in I_k, k \in K$. Dann gilt:

- (i) $\times_{k \in K} \bigcup_{i \in I_k} A_{k,i} = \bigcup_{f \in \times_{k \in K} I_k} \times_{k \in K} A_{k,f(k)},$
- (ii) $\prod_{k \in K} \sum_{i \in I_k} \alpha_{k,i} = \sum_{f \in \times_{k \in K} I_k} \prod_{k \in K} \alpha_{k,f(k)},$
- (iii) $\times_{k \in K} \bigcap_{i \in I_k} A_{k,i} = \bigcap_{f \in \times_{k \in K} I_k} \times_{k \in K} A_{k,f(k)},$
- (iv) $\bigcap_{k \in K} \bigcup_{i \in I_k} A_{k,i} = \bigcup_{f \in \times_{k \in K} I_k} \bigcap_{k \in K} A_{k,f(k)},$
- (v) $\bigcup_{k \in K} \bigcap_{i \in I_k} A_{k,i} = \bigcap_{f \in \times_{k \in K} I_k} \bigcup_{k \in K} A_{k,f(k)}.$

Das sieht ein bißchen abschreckend aus, und deswegen ist vielleicht ein Gedankenexperiment für die ersten beiden Gleichungen vertrauenerweckend und hilfreich: Jedes $k \in K$ ist ein Land, I_k sind die Städte im Land k , $A_{k,i}$ sind die Einwohner der Stadt i im Land k . Für die Produktbildung betrachten wir alle möglichen Menschenketten (mit Kerzen, wenn Sie wollen), bei denen jedes Land genau einen Bewohner aus irgendeiner seiner Städte beiträgt. All diese Ketten können wir gruppieren nach Ketten f aus Städten, bei denen jedes Land k eine Stadt $f(k)$ beiträgt, was zur rechten Summe führt. Das Distributivgesetz (i) für Mengen stimmt auch im nichtdisjunkten Fall, bei dem jeder Bewohner Wohnsitze in mehreren Städten haben kann.

Wir verwenden zuweilen folgende suggestive Schreibweisen: $\sum_{i \in I} b$ ist nichts als $\sum_{i \in I} \alpha_i$ mit $\alpha_i = b$ für alle $i \in I$. Damit ist zum Beispiel $\sum_{i \in I} 1$ definiert. Weiter ist für eine Menge $\mathfrak{A}$ von Kardinalzahlen $\sum_{\alpha \in \mathfrak{A}} \alpha$ oder auch nur $\sum \mathfrak{A}$ einfach eine bequeme Schreibweise für $\sum_{i \in \mathfrak{A}} b_i$ mit $b_i = i$ für alle $i \in \mathfrak{A}$. Analoges gilt für das Kreuzprodukt und das Mengenprodukt; so ist etwa

$$\times A = \times_{i \in A} i = \{ f \mid f : A \rightarrow \bigcup A, f(a) \in a \text{ für alle } a \in A \}.$$

In der Literatur wird das große griechische Pi oft auch für das Mengenprodukt verwendet; weiter findet man bei Zermelo und Hausdorff auch $\mathfrak{P}$ für $\prod$ bzw. $\times$.

Die Idee, daß Multiplikation iterierte Summation, und Exponentiation iterierte Multiplikation ist, behandelt die folgende Übung.

Übung

Sei A eine Menge, und sei $\alpha = |A|$. Dann gilt:

- (i) $\alpha = \sum_{i \in A} 1$,
- (ii) $b \cdot \alpha = \sum_{i \in A} b$ für alle Kardinalzahlen b ,
- (iii) $2^\alpha = \prod_{i \in A} 2$, und allgemeiner
- (iv) $b^\alpha = \prod_{i \in A} b$ für alle Kardinalzahlen b .

Ist $J \subseteq I$ unendlich, $\alpha_i \geq 2$ für alle $i \in J$, $\alpha_i \geq 1$ für alle $i \in I$, so ist nach den beiden Übungen $\prod_{i \in I} \alpha_i \geq \prod_{i \in J} 2 \geq \prod_{i \in \mathbb{N}} 2 = 2^{\aleph_0}$. Alle nichttrivialen „echt“ unendlichen Produkte haben also mindestens die Mächtigkeit der reellen Zahlen. Es gibt also i. a. sehr viele Transversalfunktionen. Andererseits gilt $\prod_{i \in I} \alpha_i \leq b^{|I|}$, falls $\alpha_i \leq b$ für alle $i \in I$ gilt.

Nach dieser Seelandschaft mit Frakturbuchstaben kommen wir nun endlich zu einem interessanten Satz.

Der Satz von Julius König und Ernst Zermelo

Den Satz von Cantor kann man nun etwas schrullig so notieren:

$$\sum_{i \in A} 1 < \prod_{i \in A} 2 \quad \text{für alle Mengen } A.$$

[Für $A = \emptyset$ ist die linke Seite 0, die rechte ist 1 wegen $\emptyset : \emptyset \rightarrow \{0, 1\}$.]

Es gilt nun die allgemeinste denkbare Form dieses strikten Größenunterschiedes zwischen Summe und Produkt. Dies ist der Inhalt eines der stärksten Sätze

der elementaren Kardinalzahlarithmetik. Der Beweis ist, wie kaum anders zu erwarten, ein Diagonalargument.

Satz (*Satz von Julius König und Ernst Zermelo*)

Sei I eine Menge, und seien α_i, b_i Kardinalzahlen für $i \in I$.

Weiter gelte $\alpha_i < b_i$ für alle $i \in I$. Dann gilt:

$$\sum_{i \in I} \alpha_i < \prod_{i \in I} b_i.$$

Beweis

Offenbar $\sum_{i \in I} \alpha_i \leq \prod_{i \in I} b_i$.

Seien A_i, B_i Mengen mit $|A_i| = \alpha_i, |B_i| = b_i$ für $i \in I$.

Sei weiter $S = \bigcup_{i \in I} A_i \times \{i\}$, $P = \prod_{i \in I} B_i$.

Sei $F: S \rightarrow P$. Wir zeigen, daß F nicht surjektiv ist.

Wir definieren $g \in P$ wie folgt. Für $i \in I$ sei:

$$g(i) = \text{„ein } b \in B_i \text{ mit } b \notin \{F(y)(i) \mid y \in A_i \times \{i\}\} \text{“}.$$

Ein solches b existiert, denn es gilt:

$$|\{F(y)(i) \mid y \in A_i \times \{i\}\}| \leq |\{F(y) \mid y \in A_i \times \{i\}\}| \leq |A_i| < |B_i|,$$

und somit ist $\{F(y)(i) \mid y \in A_i \times \{i\}\}$ eine echte Teilmenge von B_i .

Offenbar ist $g \in P$.

Aber $g \notin \text{rng}(F)$, denn $g(i) \neq F(y)(i)$ für alle $y \in A_i \times \{i\}, i \in I$,

– also $g \neq F(y)$ für alle $y \in S$.

Zermelo (1908): „33_{VI}. Theorem. Sind zwei äquivalente Menge T und T' , deren Elemente $M, N, R \dots$ bzw. $M', N', R', \dots$ unter sich elementenfremde Mengen sind, so aufeinander [bijektiv] abgebildet, daß jedes Element M von T von kleinerer Mächtigkeit ist als das entsprechende Element M' von T' , so ist auch die Summe $S = \mathfrak{C}T$ [= $\bigcup T$] aller Elemente von T von kleinerer Mächtigkeit als das Produkt $P' = \mathfrak{P}T'$ [= $\times T' = \times_{i \in T'} i$] aller Elemente von T' ...

Das vorstehende (Ende 1904 der Göttinger Mathematischen Gesellschaft von mir mitgeteilte) Theorem ist der allgemeinste bisher bekannte Satz über das Größer und Kleiner der Mächtigkeiten, aus dem alle übrigen sich ableiten lassen. Der Beweis beruht auf einer Verallgemeinerung des von Herrn J. König für einen speziellen Fall [für abzählbare Mengen T, T'] ... angewandten Verfahrens.“

Der Index VI bedeutet hier, daß Auswahlakte in den Beweis des Satzes einfließen.

Als Korollar erhalten wir den Satz von Cantor über die Mächtigkeit der Potenzmenge. Allerdings ergibt sich kein neuer Beweis für $\alpha < 2^\alpha$: Der Beweis des Satzes von König-Zermelo fällt im Fall $\alpha_i = 1$ und $b_i = 2$ für alle $i \in I$ mit dem originalen Beweis des Satzes von Cantor zusammen.

Einige Fragen zur Kardinalzahlarithmetik

Wir versammeln einige natürliche Fragen zur Kardinalzahlarithmetik, die wir noch nicht beantwortet haben. Die einfachsten sind vielleicht:

- (1) Gilt $\alpha + \alpha = \alpha$ für alle unendlichen Kardinalzahlen α ?
- (2) Gilt $\alpha \cdot \alpha = \alpha$ für alle unendlichen Kardinalzahlen α ?

Motiviert sind diese Fragen durch die Gleichungen $\alpha + \alpha = \alpha \cdot \alpha = \alpha$ für $\alpha = \aleph_0$ oder $\alpha = \mathfrak{c} = |\mathbb{R}|$ oder $\alpha = \mathfrak{f} = |\mathfrak{F}|$. Wir werden sie unten positiv beantworten.

Neben der Addition und der Multiplikation kann man Fragen an die Exponentiation stellen, insbesondere an 2^α . Weiter gibt es jede Menge Fragen an die Struktur der Ordnung der Kardinalzahlen. Wir haben ihre Vergleichbarkeit gezeigt, aber das legt die Ordnung sicher noch nicht fest.

Zur bequemen Formulierung derartiger Fragestellungen brauchen wir noch zwei sich sehr ähnliche Definitionen.

Definition (*kardinaler Nachfolger, α^+*)

Sei α eine Kardinalzahl, und sei b eine Kardinalzahl mit:

- (i) $\alpha < b$.
- (ii) Für alle Kardinalzahlen c mit $\alpha < c$ ist $b \leq c$.

Dann heißt b der *kardinale Nachfolger* von α , in Zeichen $b = \alpha^+$.

Definition (*kardinales Supremum, $\sup(\mathfrak{A})$*)

Sei $\mathfrak{A}$ eine Menge von Kardinalzahlen, und sei b eine Kardinalzahl mit:

- (i) $\alpha \leq b$ für alle $\alpha \in \mathfrak{A}$.
- (ii) Ist c eine Kardinalzahl mit $\alpha \leq c$ für alle $\alpha \in \mathfrak{A}$, so ist $b \leq c$.

Dann heißt b das (*kardinale*) *Supremum* von $\mathfrak{A}$, in Zeichen $b = \sup(\mathfrak{A})$.

Hier nun also zwölf weitere Fragen:

- (I) Existiert α^+ für alle α ?
- (II) Existiert $\sup(\mathfrak{A})$ für jede Menge $\mathfrak{A}$ von Kardinalzahlen?
- (III) Gibt es Kardinalzahlen α_n , $n \in \mathbb{N}$, mit $\alpha_{n+1} < \alpha_n$ für alle n ?
- (IV) Gilt $\sup(\mathfrak{A}) = \sum_{\alpha \in \mathfrak{A}} \alpha$ für jede Menge $\mathfrak{A}$ von Kardinalzahlen?
- (V) Gibt es ein $b > \aleph_0$, das kein Nachfolger eines α ist?
- (VI) Gibt es ein $b > \aleph_0$ mit $2^\alpha < b$ für alle $\alpha < b$?
- (VII) Gilt $2^\alpha = \alpha^+$ für ein unendliches α , falls α^+ existiert?
- (VIII) Gibt es ein unendliches α mit $2^\alpha = \alpha^+$?
- (IX) Gilt $2^\alpha < 2^b$ für gewisse (oder alle) $\alpha < b$?
- (X) Welche Werte außer $b \leq \alpha$ kann man für 2^α ausschließen?
- (XI) Gilt $|\{b \mid b < \alpha\}| \leq \alpha$ für alle Kardinalzahlen α ?
- (XII) Existiert für alle Mengen A mit $|A| = \alpha$ eine $\subseteq$ -Kette K in A mit $\{|X| \mid X \in K\} = \{b \mid b < \alpha\}$?

Übung

Es sind äquivalent:

- (i) Die Antworten auf Fragen (I) und (II) sind „ja“.
- (ii) Die Antwort auf Frage (III) ist „nein“.

[Wir geben unten einen Beweis von $(ii) \hookrightarrow (i)$.]

Viele dieser Fragen sind dadurch motiviert, daß die endlichen Kardinalzahlen und $\aleph_0$ bestimmte Eigenschaften haben; es existiert für $n \in \mathbb{N}$ z. B. immer der kardinale Nachfolger $n^+ = n + 1$ und es gilt $\aleph_0 = \sup(\mathbb{N}) = \sum_{n \in \mathbb{N}} n$. Weiter ist $\aleph_0$ kein Nachfolger und sogar abgeschlossen unter der Exponentiation 2^α , d. h. es gilt $2^\alpha < \aleph_0$ für alle $\alpha < \aleph_0$. Andere Einträge in der Liste sind natürliche Fragen an die Potenzmengenoperation.

Man kann über die Liste ein wenig nachdenken, und dann einfach in die Sprechstunde der alten Dame Mengenlehre gehen und sie fragen:

Wie sieht die Struktur der Ordnung der Kardinalzahlen aus?

Was kann man zur Exponentiation sagen, außer daß $2^\alpha > \alpha$ gilt?

Es zeigt sich, daß Madame $\mathcal{M}$ recht bereitwillig auf Strukturfragen Antwort gibt, und wir werden ihre Antworten gleich und dann noch einmal im zweiten Abschnitt besprechen. Besonders die ordnungstheoretischen Methoden ergeben ein kristallklares Bild über die Lage der Kardinalzahlen untereinander und betten sie in einen feineren Rahmen ein.

Zu Fragen der Exponentiation hüllt sich die Dame derart in Schweigen, als hätte man etwas Unanständiges gefragt. Es braucht wahrlich einen listenreichen Odysseus, um hier Antworten zu erhalten. Einiges läßt sich über 2^α herausfinden, aber oft nur mit sehr komplexen Techniken.

Viele Fragen über 2^α sind aber zumindest hart an der Grenze der Erforschbarkeit, wenn nicht sogar hinter ihr. Man hat zeigen können, daß die Beweisluft in exponentialen Höhen sehr dünn wird. (Vgl. die Diskussion über (CH) im vorherigen Abschnitt.) Manche haben in ihrer Ratlosigkeit der Dame vorgeworfen, sie wisse selber nicht so recht, was es mit 2^α auf sich habe. So weit wollen wir nicht gehen, und Fragen auch dann als sinnvoll erachten, wenn es noch unzählige Wellen brauchen wird, um von ihnen auch nur ein Körnchen Erkenntnis abzuspalten.

Wir wenden uns nun zunächst der Addition und Multiplikation von unendlichen Kardinalzahlen zu: Wir zeigen, daß diese Operationen viel einfacher sind als die Addition und Multiplikation, die wir in der Grundschule für die endlichen Nußhaufen durch Nachkochen von komplizierten Rezepten gelernt haben. Das Ergebnis ist hübsch, aber angesichts der Unantastbarkeit der Exponentiation kann man es auch als eine Art Hohngelächter empfinden.

Danach bringen wir ein Beispiel für ein Nichtsdestotrotz-Resultat über 2^α : Wir können tatsächlich einige Werte ausschließen.

Zum Abschluß des Kapitels erforschen wir dann noch recht detailliert die Strukturfragen an die Ordnung, und haben dann insgesamt viele Fragen der Liste beantwortet.

Addition und Multiplikation von Kardinalzahlen

Wir zeigen $\alpha + \alpha = \alpha \cdot \alpha = \alpha$ für alle unendlichen Kardinalzahlen α .

Die Beweisidee ist die des Vergleichbarkeitssatzes. Wie dort sind die dahinterliegenden Ideen im Grunde sehr anschaulich: Wir definieren Approximationen an ein gesuchtes Objekt und zeigen, daß es eine bestmögliche Approximation gibt. Wir benutzen hierbei den allgemeinen „des Pudels Kern“-Satz über die Existenz von Zielen in Zermelosystemen, sodaß der Leser, der nicht gerne in die Rolle des Pudels schlüpft, die etwas filzige Approximationstechnik selber nicht zu kennen braucht. (Es ist eine gute Übung, einen der Beweise unten umständlich auszuführen, und sich das mephistophelische Schauspiel von Schmidtscher Auswahl und Zermeloscher Reduzierung von geschlossenen Mengen noch einmal vor Augen zu führen.) Der Leser kann den Rest dieses Abschnitts auch überschlagen: Wir geben im zweiten Abschnitt Beweise von aufregender Transparenz für die Addition und Multiplikation von Kardinalzahlen.

Wir führen den Beweis in der „neuen“ Form mit Kardinalzahlen ($|A| = \alpha$). Ein Umschreiben auf die „alte“ relationale Form ($|A| = |B|$) wäre problemlos. Die Multiplikationsaussage lautet z. B. einfach: Für alle unendlichen Mengen M ist $|M \times M| = |M|$.

Der Beweis gliedert sich in mehrere Zwischenstufen. Zunächst zeigen wir, daß wir eine unendliche Menge immer als „Rechteck“ mit einer abzählbaren Seite darstellen können.

Satz (Zerlegungslemma)

Sei A eine unendliche Menge.

Dann gibt es eine Menge B mit $|A| = |B \times \mathbb{N}|$.

Die Idee ist, A durch abzählbare Teilmengen auszuschöpfen, um es in der Form $B \times \mathbb{N}$ anordnen zu können. Jede dabei verwendete Teilmenge bildet dann eine Spalte der Höhe $\mathbb{N}$ in der abschließenden Darstellung als „Rechteck“. Die ersten Elemente der Spalten bilden die Breite des Rechtecks, also die Menge B .

Beweis

Sei $\mathcal{Z} = \{ Z \mid \text{jedes } f \in Z \text{ ist eine Injektion von } \mathbb{N} \text{ nach } A \text{ und} \\ \text{für alle } f, g \in Z \text{ mit } f \neq g \text{ gilt: } \text{rng}(f) \cap \text{rng}(g) = \emptyset \}.$

Dann ist $\mathcal{Z}$ ein Zermelosystem!

Wegen $|\mathbb{N}| \leq |A|$ ist $\mathcal{Z} \neq \emptyset$.

Sei also Z ein beliebiges Ziel von $\mathcal{Z}$. Dann gilt:

($\diamond$) $R = A - \bigcup_{f \in Z} \text{rng}(f)$ ist endlich.

Hieraus folgt aber:

($\diamond^*$) Es existiert ein $Z^* \in \mathcal{Z}$ mit $\bigcup_{f \in Z^*} \text{rng}(f) = A$.

Beweis von ($\diamond^$)*

Sei $h: \bar{m} \rightarrow R$ bijektiv für ein $m \in \mathbb{N}$, und sei $g \in Z$ beliebig.

Wir definieren $g^* : \mathbb{N} \rightarrow A$ durch

$$g^*(n) = \begin{cases} h(n), & \text{falls } n < m, \\ g(n-m), & \text{andernfalls.} \end{cases}$$

Dann ist g^* injektiv und $\text{rng}(g^*) = \text{rng}(g) \cup R$.

Sei $Z^* = (Z - \{g\}) \cup \{g^*\}$.

Dann ist $Z^* \in \mathcal{X}$ und $\bigcup_{f \in Z^*} \text{rng}(f) = A$.

Sei nun Z^* wie in $(\diamond^*)$. Sei $B = \{f(0) \mid f \in Z^*\}$.

Wir definieren $h : B \times \mathbb{N} \rightarrow A$ durch:

$h(x, n) = f_x(n)$, wobei f_x das eindeutige $f \in Z^*$ ist mit $f(0) = x$.

- Dann ist $h : B \times \mathbb{N} \rightarrow A$ bijektiv.

Mit einem kleinen Trick erhalten wir hieraus schon:

Korollar (*abzählbarer Multiplikationssatz*)

Sei α eine unendliche Kardinalzahl.

Dann gilt $\alpha \cdot \aleph_0 = \alpha$.

Beweis

Sei A eine Menge mit $|A| = \alpha$. Nach dem Zerlegungslemma existiert eine Menge B mit $|A| = |B \times \mathbb{N}|$. Sei $b = |B|$.

Wir haben damit $\alpha = b \cdot \aleph_0$. Dann gilt wegen $\aleph_0 \cdot \aleph_0 = \aleph_0$ und der Assoziativität der Multiplikation:

$$\alpha = b \cdot \aleph_0 = b \cdot \aleph_0 \cdot \aleph_0 = \alpha \cdot \aleph_0.$$

In der Rechnung der letzten Zeile des Beweises zeigt sich, warum wir im Zerlegungslemma die Menge A nicht durch Paare, sondern durch abzählbare Mengen ausgeschöpft haben: Eine Paar ausschöpfung – die zudem ein lästiges einzelnes Element hinterlassen könnte, das sich nicht ganz so leicht einbinden ließe wie ein endlicher Rest in eine $\mathbb{N}$ -Ausschöpfung – liefert nur $\alpha \cdot 2 = b \cdot 2 \cdot 2 = b \cdot 4$, was nicht weiterhilft. $2 \cdot 2 \neq 2$, aber $\aleph_0 \cdot \aleph_0 = \aleph_0$.

Wegen $\alpha \leq \alpha + \alpha = \alpha \cdot 2 \leq \alpha \cdot \aleph_0 = \alpha$ für unendliche α haben wir:

Korollar (*Additionssatz*)

Sei α eine unendliche Kardinalzahl. Dann gilt $\alpha + \alpha = \alpha$.

Mit Hilfe des Additionssatzes beweisen wir nun den Multiplikationssatz, wieder durch Approximation. Hierzu einige Bemerkungen vorab. Sei B eine Approximation an $|A \times A| = |A|$, d.h. $B \subseteq A$ und $|B \times B| = |B|$.

Ist $|B| < |A|$, so können wir im Rest $A - B$ eine Kopie C von B finden, denn „ B ist klein in A “. Die Kopie fügen wir zu B hinzu und betrachten das Kreuzprodukt von $B \cup C$. Dieses zerfällt in vier gleichgroße Teile, und mit Hilfe des Additionssatzes können wir leicht zeigen, daß wir einen Zeugen f für $|B \times B| = |B|$ zu einem Zeugen g für $|B \cup C \times B \cup C| = |B \cup C|$ fortsetzen können. Die Approximation B ist also keinesfalls bestmöglich, und dies hält die Dinge intern am Laufen.

Ist *andernfalls* $|B| = |A|$, so ...? – ja, so sind wir schon fertig, denn in diesem Fall ist $|A \times A| = |A|$, da dies für B gilt, und B und A in bezug auf Kardinalitätsfragen gleichwertig sind. Dies ist eine kleine, aber enorm hilfreiche Beobachtung: Wir müssen nicht ganz A ausschöpfen, es reicht, daß wir eine Approximation $B \subseteq A$ finden mit $|B \times B| = |B|$, die die gleiche Kardinalität wie A hat.

Satz (*Multiplikationssatz*)

Sei α eine unendliche Kardinalzahl.
Dann gilt $\alpha \cdot \alpha = \alpha$.

Beweis

Sei A eine Menge mit $|A| = \alpha$.
Sei $\mathcal{L} = \{ f \mid f : B \times B \rightarrow B \text{ bijektiv für ein unendliches } B \subseteq A \}$.
Es gilt $\mathcal{L} \neq \emptyset$ wegen $\aleph_0 \leq \alpha$ und $\aleph_0 \cdot \aleph_0 = \aleph_0$.
 $\mathcal{L}$ ist ein Zermelosystem.
Sei also $f \in \mathcal{L}$ ein Ziel von $\mathcal{L}$.
Sei $B = \text{rng}(f)$. Dann gilt $|B \times B| = |B|$ wegen $f \in \mathcal{L}$.
Ist $|B| = |A|$, so gilt $|A \times A| = |B \times B| = |B| = |A|$, und wir sind fertig.
Es genügt also, den anderen Fall zum Widerspruch zu führen.

Annahme, es gilt $|B| < |A|$.

Sei dann $C \subseteq A - B$ mit $|C| = |B|$.

[Ein solches C existiert, denn *andernfalls* ist $|A - \text{rng}(f)| < |B|$.

Aber $|B| < |A|$, also gilt nach dem Additionssatz

$$|A| = |A - B| + |B| \leq |B| + |B| = |B| < |A|, \text{ Widerspruch.}]$$

Sei $D = B \cup C$, $D_1 = B \times C$,

$D_2 = C \times C$, $D_3 = C \times B$,

$f^* = \text{„ein } h \supseteq f \text{ mit}$

$h : D \times D \rightarrow D \text{ bijektiv“}$.

Ein solches h existiert, denn es gilt

$|D_1| = |D_2| = |D_3| = |B|$, also

hat nach dem Additionssatz auch

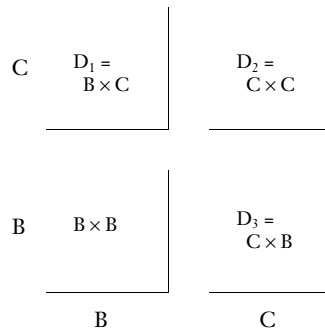
$D_1 \cup D_2 \cup D_3$ die Kardinalität

$|B| = |C|$. Also existiert ein

$g : D_1 \cup D_2 \cup D_3 \rightarrow C$ bijektiv.

Also ist $f^* = f \cup g : D \times D \rightarrow D$ bijektiv. Also $f^* \in \mathcal{L}$.

– Dann ist aber f kein Ziel von $\mathcal{L}$, *Widerspruch*.



Der hier dargestellte Beweis folgt im wesentlichen einem Beweis von Max Zorn aus dem Jahre 1944. Die ersten Beweise des Additionssatzes und des Multiplikationssatzes haben unabhängig A. E. Harward 1905 und Gerhard Hessenberg 1906 gegeben, zuvor war $\alpha = \alpha \cdot \alpha = \alpha + \alpha$ nur für die Spezialfälle $\aleph_0$, $\aleph_1$ und $\aleph_2$ bekannt. Bernstein schreibt in seiner Doktorarbeit 1901, daß ihm Cantor einen Beweis der Gleichung $|M \times M| = |M|$ mitgeteilt habe für Mengen M mit einer bestimmten Eigenschaft (E), gibt diesen Beweis aber nicht wieder. Hausdorff stellt 1904 eine äquivalente Behauptung auf, ebenfalls

ohne Beweis. Die damals kursierende und naheliegende Beweisidee war die der Verallgemeinerung der Cantorschen Diagonalaufzählung von $\mathbb{N}^2$. Mengen M mit der Eigenschaft (E) lassen sich in eine $\mathbb{N}$ -ähnliche Form bringen, und es liegt dann nahe, das Produkt $M \times M$ über eine Paarungsfunktion diagonal abzuzählen.

Harward und Hessenberg gelangen die ersten strengen Umsetzungen dieser Idee. Ein Jahr später gab Hessenberg einen zweiten Beweis, und 1908 fand Philip Jourdain (1879 – 1921) einen dritten. Wir werden diese Beweise und ihre Geschichte später diskutieren (2.8). Sie alle zeigen den Multiplikationssatz für Mengen, die die Eigenschaft (E) haben. Das allgemeine Resultat folgt dann aus dem Satz von Zermelo, daß in der Tat alle Mengen die Eigenschaft (E) haben. Eigenschaft (E) ist die sog. Wohlordenbarkeit, und der Satz von Zermelo ist der Wohlordnungssatz von 1904, ein historischer Vulkanausbruch, der viel fruchtbaren Boden hinterließ. Der Durchführung des Cantorschen Leitmotivs *Wohlordnung* widmet sich der zweite Abschnitt, und wir werden auf die Ereignisse kurz nach der Jahrhundertwende dort zurückkommen.

Als Korollar halten wir fest, daß die Addition und Multiplikation im Unendlichen letztendlich trivial ist:

Korollar (*Addition und Multiplikation von Kardinalzahlen*)

Seien $\alpha, b \neq 0$ Kardinalzahlen, und es sei $\aleph_0 \leq \max(\alpha, b)$.

Dann gilt $\alpha + b = \alpha \cdot b = \max(\alpha, b)$.

Hatten wir auch etwas Mühe mit dem Beweis des Multiplikationssatzes, so bleibt uns dafür die Mühe des Ausrechnens von $\alpha \cdot b$ für immer erspart.

Aus dem Multiplikationssatz gewinnen wir die folgende Aussage über die Kardinalzahlsumme und über Suprema: Sei $\mathfrak{A}$ eine Menge von unendlichen Kardinalzahlen. Für alle Kardinalzahlen c gelte $|\{b \mid b < c\}| \leq c$. Dann gilt:

$$\sup(\mathfrak{A}) = \sum_{\alpha \in \mathfrak{A}} \alpha.$$

Denn sei c eine Kardinalzahl mit $\alpha \leq c$ für alle $\alpha \in \mathfrak{A}$. Dann gilt:

$$\sum_{\alpha \in \mathfrak{A}} \alpha \leq c \cdot |\mathfrak{A}| \leq c \cdot c = c.$$

Offenbar ist $\alpha \leq \sum_{\alpha \in \mathfrak{A}} \alpha$ für alle $\alpha \in \mathfrak{A}$, und damit ist dann $\sup(\mathfrak{A}) = \sum_{\alpha \in \mathfrak{A}} \alpha$. Wir zeigen unten, daß die Voraussetzung $|\{b \mid b < c\}| \leq c$ für Kardinalzahlen c immer erfüllt ist.

Früher hatten wir gezeigt, daß das Entfernen einer abzählbaren Menge die Kardinalität einer unendlichen Menge nicht ändert. Allgemein erhalten wir nun:

Korollar (*Subtraktionssatz*)

Sei A eine unendliche Menge, und sei $B \subseteq A$ mit $|B| < |A|$.

Dann gilt $|A - B| = |A|$.

Beweis

Andernfalls wäre $|A| = |B| + |A - B| = \max(|B|, |A - B|) < |A|$,

– *Widerspruch.*

Eine hübsche Anwendung des Multiplikationssatzes ist, daß wir α^b in vielen Fällen auf 2^b zurückführen können:

Übung (Verkleinern der Basis)

Seien α, b Kardinalzahlen mit $2 \leq \alpha \leq b$, $\aleph_0 \leq b$.

Dann gilt $\alpha^b = 2^b$.

$$[\alpha^b \leq (2^\alpha)^b = 2^{\alpha \cdot b} = 2^b.]$$

Wir betrachten nun umgekehrt den Fall, daß der Exponent kleiner ist als die Basis. Hier gibt es keine einfache Formel, dafür aber einen interessanten Zusammenhang zwischen α^b und der Anzahl der b -großen Teilmengen einer α -großen Menge. Hierzu:

Definition ($[A]^b, [A]^{\leq b}, [A]^{< b}$)

Sei A eine Menge, und sei $b \leq |A|$. Dann setzen wir:

$$[A]^b = \{ B \subseteq A \mid |B| = b \},$$

$$[A]^{\leq b} = \{ B \subseteq A \mid |B| \leq b \},$$

$$[A]^{< b} = \{ B \subseteq A \mid |B| < b \}.$$

Sind A, B Mengen, so gilt ${}^B A \subseteq [B \times A]^b$, denn jedes $f: B \rightarrow A$ ist eine Teilmenge von $B \times A$ der Mächtigkeit $b = |B|$. Diese Beobachtung ist schon fast die Hälfte des folgenden Satzes:

Satz (α^b und $[A]^b$)

Seien A eine unendliche Menge, $\alpha = |A|$, und sei $b \leq \alpha$ eine Kardinalzahl.

Dann gilt $\alpha^b = |[A]^b| = |[A]^{\leq b}|$.

Beweis

Sei B eine Menge mit $|B| = b$.

O. E. $b \neq 0$, denn sonst ist ${}^B A = [A]^b = [A]^{\leq b} = \{\emptyset\}$.

zu $\alpha^b \leq |[A]^b|$:

Es gilt $\alpha^b = |{}^B A| \leq |[B \times A]^b| = |[A]^b|$ wegen $|B \times A| = |A|$.

zu $|[A]^b| \leq |[A]^{\leq b}|$:

Es gilt $[A]^b \subseteq [A]^{\leq b}$.

zu $|[A]^{\leq b}| \leq \alpha^b$:

Wir definieren $h(C)$ für alle $C \subseteq A$ mit $|C| \leq b$, $C \neq \emptyset$, durch

$h(C) = \text{„ein surjektives } f: B \rightarrow C\text{“}$.

Dann ist $h: [A]^{\leq b} - \{\emptyset\} \rightarrow {}^B A$ injektiv.

— Also (wegen $b \neq 0$) $|[A]^{\leq b}| = |[A]^{\leq b} - \{\emptyset\}| \leq |{}^B A|$.

Zerlegungen und ein Resultat über $|\mathbb{R}|$

Wir haben gesehen, daß wir eine unendliche Kardinalzahl nicht als Summe zweier kleinerer Kardinalzahlen darstellen können. Sicher können wir aber α als Summe von α -vielen kleinen Mengen darstellen, etwa trivial als $\alpha = \sum_{i \in I} 1$ mit einer Menge I mit $|I| = \alpha$. Eine natürliche Frage ist nun die Untersuchung von Zwischenstufen, bei denen α als Summe von weniger als α -vielen Kardinalitäten, die alle kleiner als α sind, dargestellt wird.

Definition (*c*-zerlegbar)

Seien α, c Kardinalzahlen. α heißt *c*-zerlegbar, falls gilt:

Es gibt eine Menge I und Kardinalzahlen b_i für $i \in I$ mit:

- (i) $|I| = c$,
- (ii) $b_i < \alpha$ für alle $i \in I$,
- (iii) $\sum_{i \in I} b_i = \alpha$.

Wir sagen dann, daß $\langle b_i \mid i \in I \rangle$ eine *c*-Zerlegung von α ist.

Jede Kardinalzahl $\alpha \neq 1$ ist trivialerweise α -zerlegbar. Wir können aber leicht recht große Kardinalzahlen konstruieren, die $\aleph_0$ -zerlegbar sind:

Übung

Sei α_0 eine Kardinalzahl, und sei A_0 eine Menge mit $|A_0| = \alpha_0$.

Wir definieren rekursiv für $n \in \mathbb{N}$: $A_{n+1} = \mathcal{P}(A_n)$.

Sei $A = \bigcup_{n \in \mathbb{N}} A_n$, und sei $\alpha = |A|$.

Dann gilt $\alpha > \alpha_0$ und α ist $\aleph_0$ -zerlegbar.

Wir können überraschenderweise nun alle $\aleph_0$ -zerlegbaren Kardinalzahlen als mögliche Werte der Kardinalität des Kontinuums ausschließen, und also doch ein bißchen mehr über $2^{\aleph_0}$ herausfinden als nur $2^{\aleph_0} > \aleph_0$!

Satz (*$\aleph_0$ -Unzerlegbarkeit von $|\mathbb{R}|$*)

Sei α eine unendliche Kardinalzahl, und α sei $\aleph_0$ -zerlegbar.

Dann gilt $|\mathbb{R}| \neq \alpha$.

Beweis

Sei $c = 2^{\aleph_0}$, und seien b_n Kardinalzahlen mit $b_n < c$ für alle $n \in \mathbb{N}$.

Dann gilt nach dem Satz von König-Zermelo:

$$\sum_{n \in \mathbb{N}} b_n < \prod_{n \in \mathbb{N}} c = |\mathbb{N}\mathbb{R}| = c^{\aleph_0} = (2^{\aleph_0})^{\aleph_0} = 2^{\aleph_0 \cdot \aleph_0} = 2^{\aleph_0} = c.$$

— Also ist $\langle b_n \mid n \in \mathbb{N} \rangle$ keine $\aleph_0$ -Zerlegung von c .

Allgemeiner zeigt das Argument:

Übung

Sei α eine unendliche Kardinalzahl.

Dann ist 2^α nicht c -zerlegbar für alle $c \leq \alpha$.

Es gibt also sowohl beliebig große $\aleph_0$ -zerlegbare Kardinalzahlen, als auch solche, die nicht α -zerlegbar sind für ein beliebig großes α .

Zerlegungen wurden von Hessenberg und Hausdorff im ersten Jahrzehnt des 20. Jahrhunderts untersucht, wobei sie dort in ordnungstheoretischem Gewande auftreten. Wir kommen bei der Besprechung des Konfinalitätsbegriffs darauf noch zurück. Zerlegungen bzw. Konfinalitäten sind von großer Bedeutung und tauchen in der modernen Mengenlehre an allen Ecken und Enden auf. Für jetzt genügt es uns, etwas über $|\mathbb{R}|$ gefunden zu haben, was wir bislang nicht wußten. Dem interessierten Leser sei aber noch eine Übung angeboten.

Übung

Sei α eine unendliche Kardinalzahl, und es existiere α^+ .

Dann ist α^+ c -unzerlegbar für alle $c < \alpha^+$.

[Benutze $b \leq \alpha$ für alle Kardinalzahlen $b < \alpha^+$ und $\alpha \cdot \alpha = \alpha$.]

Zur Struktur der Ordnung der Kardinalzahlen

Den Nachweis der Existenz des Supremums einer Menge $\mathfrak{A}$ von Kardinalzahlen sind wir noch schuldig geblieben. Wir betrachten zunächst Ketten.

Satz (*Suprema und Vereinigungen von Ketten*)

Sei K eine $\subseteq$ -Kette, und sei $B = \bigcup K$.

Weiter seien $\mathfrak{A} = \{ |X| \mid X \in K \}$, $b = |B|$, und es sei c eine Kardinalzahl mit $\alpha < c$ für alle $\alpha \in \mathfrak{A}$.

Dann gilt $b \leq c$.

$$b = |B| = \left| \bigcup K \right|$$

$$c = |C|$$

⋮

Beweis

Die Aussage ist klar, falls c endlich ist. O. E. sei also c eine unendliche Kardinalzahl.

Sei $K' = \{ \bigcup U \mid U \subseteq K \}$.

Dann ist K' eine Kette und es gilt:

(i) $\bigcup K' = B$,

(ii) für alle $X \in K'$, $X \neq B$, existiert ein $Y \in K$ mit $X \subset Y$,

(iii) $\bigcup U \in K'$ für alle $U \subseteq K'$.

Sei C eine Menge mit $|C| = c$.

Sei $\mathcal{Z} = \{ f \mid f : X \rightarrow C \text{ injektiv für ein } X \in K' \}$.

Dann ist $\mathcal{Z}$ ein Zermelosystem (wegen (iii)).

Sei also $f : X \rightarrow C$ ein Ziel von $\mathcal{Z}$. Dann gilt $X = B$:

⋮
 $\mathfrak{A}$ (Kardinalitäten
 der Elemente von K)
 ⋮

Wir zeigen, daß eine Kardinalzahl c zwischen $\mathfrak{A}$ und b nicht existiert.

Denn *andernfalls* existiert nach (ii) ein $Y \in K$ mit $X \subset Y$. Wegen $|\text{rng}(f)| = |X| \leq |Y| < |C|$ ist $|Y - X| \leq |Y| < |C| = |C - \text{rng}(f)|$. Also existiert ein injektives $g : Y - X \rightarrow C - \text{rng}(f)$. Dann ist aber $f \cup g : Y \rightarrow C$ injektiv. Also ist f kein Ziel von $\mathcal{L}$, *Widerspruch*.

– Dann ist aber $f : B \rightarrow C$ injektiv, also $b \leq c$.

Hat $\mathfrak{A}$ kein größtes Element, so zeigt der Beweis, daß $b = \sup(\mathfrak{A})$ gilt. Andernfalls zeigt der Beweis nur, daß b kleinergleich jedem c ist, das *strikt* größer ist als alle $a \in \mathfrak{A}$. (Sowohl $b = \sup(\mathfrak{A})$ als auch $b = \sup(\mathfrak{A})^+$ sind in diesem Fall möglich, wie später klar werden wird.)

Wir erhalten hieraus durch ein weiteres einfaches Zermelosystem-Argument, daß es auch für unendliche Kardinalzahlen kein unendliches Rückwärtszählen gibt. Der Leser, der die Übung nach den Fragen (I) – (XII) verfolgt hat, wird sich erinnern, daß dies eine sehr nützliche und starke Eigenschaft ist.

Satz (*Nichtexistenz unendlicher absteigender Folgen von Kardinalzahlen*)

Es gibt keine Kardinalzahlen α_n , $n \in \mathbb{N}$, mit $\alpha_{n+1} < \alpha_n$ für alle $n \in \mathbb{N}$.

Beweis

Wir nennen eine Kardinalzahl α (für diesen Beweis) *irreal*, falls ein kardinaler Abstieg $\alpha > \alpha_0 > \alpha_1 > \dots > \alpha_n > \dots$, $n \in \mathbb{N}$, existiert. Andernfalls nennen wir α *real*.

$\aleph_0$ ist real, denn ist $k_{n+1} < k_n$ für alle $n \in \mathbb{N}$, $k_n \in \mathbb{N}$, so hätte die nichtleere Menge $A = \{k_n \mid n \in \mathbb{N}\} \subseteq \mathbb{N}$ kein kleinstes Element.

Weiter gilt offenbar: Ist α real und $b < \alpha$, so ist auch b real.

(Denn andernfalls wäre $\alpha > b > b_0 > b_1 > b_2 \dots$ für gewisse b_n , $n \in \mathbb{N}$.)

Sei nun α eine unendliche Kardinalzahl, und sei A eine Menge mit $|A| = \alpha$.

Sei $\mathcal{L} = \{X \subseteq A \mid |X| \text{ ist real}\}$.

Dann ist $\mathcal{L}$ ein Zermelosystem:

Denn sei $K \subseteq \mathcal{L}$ eine Kette, und sei $B = \bigcup K$, $b = |B|$.

Annahme, es gibt Kardinalzahlen $b = b_0 > b_1 > b_2 > \dots > b_n > \dots$, $n \in \mathbb{N}$.

Wegen $b_1 < b_0 = b$ existiert dann nach dem Satz oben ein $X \in K$ mit $b_1 \leq |X|$. Dann ist aber $|X|$ irreal, *Widerspruch*.

Sei also $X \in \mathcal{L}$ ein Ziel von $\mathcal{L}$. Dann gilt $X = A$:

Andernfalls existiert ein $y \in A - X$.

Offenbar ist X unendlich.

Dann aber $|X| = |X \cup \{y\}|$, also $X \cup \{y\} \in \mathcal{L}$.

Also ist X kein Ziel von $\mathcal{L}$, *Widerspruch*.

– Also ist $X = A \in \mathcal{L}$, und damit ist $\alpha = |A|$ real.

Definition ($\min(\mathfrak{A})$)

Sei $\mathfrak{A}$ eine Menge von Kardinalzahlen, und sei α eine Kardinalzahl.

α heißt das *Minimum von* $\mathfrak{A}$, in Zeichen $\alpha = \min(\mathfrak{A})$, falls gilt:

$\alpha \in \mathfrak{A}$ und $\alpha \leq b$ für alle $b \in \mathfrak{A}$.

Korollar (*Existenz minimaler Elemente*)

Sei $\mathfrak{A}$ eine nichtleere Menge von Kardinalzahlen. Dann existiert $\min(\mathfrak{A})$.

Beweis

Sei $\alpha_0 \in \mathfrak{A}$ beliebig. Wir definieren durch Induktion nach $n \in \mathbb{N}$ solange möglich:

$\alpha_{n+1} =$ „ein $\alpha \in \mathfrak{A}$ mit $\alpha < \alpha_n$ “.

Nach dem Satz oben existiert ein letztes definiertes α_n ,

- und dann ist offenbar $\alpha_n = \min(\mathfrak{A})$.

Korollar (*Existenz von Nachfolgerkardinalzahlen*)

Sei α eine Kardinalzahl. Dann existiert α^+ .

Beweis

Sei $b > \alpha$ beliebig, und sei $\mathfrak{C} = \{c \mid c \text{ Kardinalzahl}, \alpha < c \leq b\}$.

- Dann gilt $\mathfrak{C} \neq \emptyset$, und es gilt $\alpha^+ = \min(\mathfrak{C})$.

Korollar (*Existenz von Suprema*)

Sei $\mathfrak{A}$ eine Menge von Kardinalzahlen. Dann existiert $\sup(\mathfrak{A})$.

Beweis

Sei b eine Kardinalzahl mit $b \geq \alpha$ für alle $\alpha \in \mathfrak{A}$.

Ein solches b existiert: Für $\alpha \in \mathfrak{A}$ sei

$A_\alpha =$ „eine Menge A mit $|A| = \alpha$ “.

Sei $B = \bigcup_{\alpha \in \mathfrak{A}} A_\alpha$. Dann ist $|B| \geq \alpha$ für alle $\alpha \in \mathfrak{A}$.

Sei $\mathfrak{C} = \{c \mid c \text{ Kardinalzahl}, c \leq b, \alpha \leq c \text{ für alle } \alpha \in \mathfrak{A}\}$.

- Dann ist $\mathfrak{C} \neq \emptyset$, und es gilt $\sup(\mathfrak{A}) = \min(\mathfrak{C})$.

Übung

Sei I eine Menge und seien $\alpha_i \geq 1$ Kardinalzahlen für $i \in I$.

Dann gilt $\sum_{i \in I} \alpha_i = |I| \cdot \sup_{i \in I} \alpha_i$.

Übung

Sei α eine Kardinalzahl. Dann gilt $|\{b \mid b < \alpha\}| \leq \alpha$.

[Annahme nicht für ein minimal gewähltes $\alpha = |A| > \aleph_0$.

Für $b < \alpha$ sei $X_b =$ „ein $X \subseteq A$ mit $|X| = b$ “.

Betrachte $g(b) =$ „ein $x \in X_{b^+} - \bigcup_{c \leq b} X_c$ “ für unendliche $b < \alpha$.]

Übung

Sei α eine Kardinalzahl, und sei A eine Menge mit $|A| = \alpha$.

Dann existiert eine Kette K in A (d.h. $K \subseteq \mathcal{P}(A)$) mit:

- (i) für alle $X, Y \in K$ mit $X \neq Y$ gilt $|X| \neq |Y|$,
- (ii) für alle $b < \alpha$ existiert ein $X \in K$ mit $|X| = b$.

Übung

Es gibt beliebig große Kardinalzahlen $\mathfrak{b}$ mit $2^\alpha < \mathfrak{b}$ für alle $\alpha < \mathfrak{b}$.

[Sei $\mathfrak{b}_0$ beliebig, und $\mathfrak{b}_{n+1} = 2^{\mathfrak{b}_n}$ für $n \in \mathbb{N}$. Betrachte $\mathfrak{b} = \sup_{n \in \mathbb{N}} \mathfrak{b}_n$.]

Der Leser wird den Kalkül mit Kardinalzahlen vielleicht schätzen gelernt haben und nun sehen, wie viel mehr in einer Ordnung stecken kann als nur die Vergleichbarkeit. Wie versprochen werden wir das Sujet im zweiten Abschnitt noch einmal verfilmen, und dann auch das Herz ansprechen.

Cantor hat den Kalkül der Exponentiation in den 90er Jahren des 19. Jahrhunderts entdeckt, und er konnte bei der Niederschrift der Potenzregeln für Kardinalzahlen seine Begeisterung nicht verbergen. Und mit dieser Begeisterung schließen wir dieses über weite Strecken etwas kühle und rechnerische Kapitel.

Georg Cantor über die allgemeinen Exponentiationsregeln

„[Hieraus folgen] die für drei beliebige Kardinalzahlen α , $\mathfrak{b}$ und $\mathfrak{c}$ gültigen Sätze...:

$$\begin{aligned} (8) \quad & \alpha^{\mathfrak{b}} \cdot \alpha^{\mathfrak{c}} = \alpha^{\mathfrak{b} + \mathfrak{c}}, \\ (9) \quad & \alpha^{\mathfrak{c}} \cdot \mathfrak{b}^{\mathfrak{c}} = (\alpha \cdot \mathfrak{b})^{\mathfrak{c}}, \\ (10) \quad & (\alpha^{\mathfrak{b}})^{\mathfrak{c}} = \alpha^{\mathfrak{b} \cdot \mathfrak{c}}. \end{aligned}$$

[kleingedruckt:] Wie inhaltreich und weittragend diese einfachen auf die Mächtigkeiten ausgedehnten Formeln sind, erkennt man an folgendem Beispiel:

Bezeichnen wir die Mächtigkeit des Linearkontinuums ... mit $\mathfrak{v}$, so überzeugt man sich leicht, daß sie sich unter anderem durch die Formel

$$(11) \quad \mathfrak{v} = 2^{\aleph_0}$$

darstellen läßt ...

Aus (11) folgt durch Quadrieren ...

$$\mathfrak{v} \cdot \mathfrak{v} = 2^{\aleph_0} \cdot 2^{\aleph_0} = 2^{\aleph_0 + \aleph_0} = 2^{\aleph_0} = \mathfrak{v}$$

und hieraus durch fortgesetzte Multiplikation mit $\mathfrak{v}$

$$(13) \quad \mathfrak{v}^{\mathfrak{v}} = \mathfrak{v},$$

wo $\mathfrak{v}$ irgendeine endliche Kardinalzahl ist [$\mathfrak{v} \in \mathbb{N}$, $\mathfrak{v} \geq 1$].

Erhebt man beide Seiten von (11) zur Potenz $\aleph_0$, so erhält man

$$\mathfrak{v}^{\aleph_0} = (2^{\aleph_0})^{\aleph_0} = 2^{\aleph_0 \cdot \aleph_0}.$$

Da aber ... $\aleph_0 \cdot \aleph_0 = \aleph_0$, so ist

$$(14) \quad \mathfrak{v}^{\aleph_0} = \mathfrak{v}.$$

Die Formeln (13) und (14) haben aber keine andere Bedeutung als diese: „Das $\mathfrak{v}$ -dimensionale sowohl, wie das $\aleph_0$ -dimensionale Kontinuum haben die Mächtigkeit des eindimensionalen Kontinuums.“ Es wird also *der ganze Inhalt* der Arbeit ... [Cantor, 1878] *mit diesen wenigen Strichen* aus den *Grundformeln des Rechnens mit Mächtigkeiten* rein algebraisch abgeleitet.“

(Georg Cantor 1895, „Beiträge zur Begründung der transfiniten Mengenlehre. I“)

13. Paradoxien der naiven Mengenlehre

Zum Schluß dieses Abschnitts und passend zur Kapitelzahl nutzen wir unsere Diagonalisierungserfahrung, um zu zeigen, daß die naive Komprehension

$$M = \{ x \mid \mathcal{E}(x) \}$$

von Objekten x mit einer gewissen Eigenschaft $\mathcal{E}(x)$ zu einer Menge M widersprüchlich ist. Die Mengenbildung und Existenzaussagen über Mengen müssen also sorgfältiger behandelt werden.

Wir betrachten zunächst das folgende Paradoxon, das auf Cantor zurückgeht.

Cantorsches Paradoxon

Sei $V = \{ x \mid x = x \}$ die Menge aller Objekte.

Wir betrachten $\mathcal{P}(V)$. Offenbar gilt $\mathcal{P}(V) \subseteq V$, denn $x \in \mathcal{P}(V)$ folgt $x = x$.

Definiere also $f: \mathcal{P}(V) \rightarrow V$ durch:

$$f(x) = x \text{ für } x \in \mathcal{P}(V).$$

Dann ist f injektiv, also haben wir:

$$|\mathcal{P}(V)| \leq |V|.$$

Aber nach dem Satz von Cantor gilt $|M| < |\mathcal{P}(M)|$ für alle Mengen M .

– Also $|V| < |\mathcal{P}(V)| \leq |V|$, *Widerspruch!*

Der gleiche Widerspruch ergibt sich, wenn wir statt der Menge V aller Objekte die Menge $V' = \{ x \mid x \text{ ist Menge} \}$ aller Mengen betrachten, denn auch hier gilt $\mathcal{P}(V') \subseteq V'$. Jede Inklusion $\mathcal{P}(M) \subseteq M$ für eine Menge M ist ein Widerspruch zum Satz von Cantor.

Bertrand Russell ist durch das Studium dieses Argumentes auf sein berühmtes eigenes Paradoxon gestoßen. Es wurde unabhängig auch von Ernst Zermelo entdeckt.

Russell-Zermelosches Paradoxon

Sei $R = \{ x \mid x \text{ ist Menge und } x \notin x \}$ die Menge aller Mengen, die nicht Element von sich selbst sind. Dann gilt für alle Mengen y :

$$(+)\ y \in R \text{ gdw } y \notin y.$$

Insbesondere gilt (+) auch für $y = R$. Dies ergibt:

$$R \in R \text{ gdw } R \notin R.$$

– *Widerspruch!*

Russell (1903, Kapitel 10, S. 102): „In terms of classes the contradiction appears even more extraordinary. A class as one may be a term of itself as many. Thus the class of all classes is a class; the class of all the terms that are not men is not a man, and so on. Do all the classes that have this property form a class? If so, is it as one a member of itself as many or not? If it is, then it is one of the classes which, as ones, are not members of themselves as many, and *vice versa*. Thus we must conclude again that the classes which as ones are not members of themselves as many do not form a class – or rather, that they do not form a class as one, for the argument cannot show that they do not form a class as many.“

Russell entdeckte seine Paradoxie Mitte 1901, wie er später erzählt. Er teilte sie Gottlob Frege (1848 – 1925) brieflich am 16. Juni 1902 mit. (Englische Übersetzungen des auf Deutsch geschriebenen Briefes von Russell und der Antwort von Frege vom 22. Juni 1902 finden sich in der Sammlung [Heijenoort 1967]. Vgl. auch [Frege 1903]). Unabhängig wurde die Paradoxie von Ernst Zermelo gefunden:

Zermelo (1908a): „Und doch hätte schon die elementare Form, welche Herr B. Russell** den mengentheoretischen Antinomien gegeben hat, sie [die Skeptizisten der Mengenlehre] überzeugen können, daß die Lösung dieser Schwierigkeiten ... in einer geeigneten Einschränkung des Mengenbegriffes zu suchen ist ...

**), 'The Principles of Mathematics', vol. I (Cambridge 1903), p. 366 – 368. Indessen hatte ich selbst diese Antinomie unabhängig von Russell gefunden, und sie schon vor 1903 u. a. Herrn Prof. Hilbert mitgeteilt.“

Die Paradoxie von Russell-Zermelo hängt eng mit Cantors Diagonalargument zusammen: $R = \{x \mid x \text{ ist Menge und } x \notin x\}$ ist nichts anderes als die Cantorsche Diagonalisierung $D = \{x \in M \mid x \notin f(x)\}$ aus dem Beweis der Ungleichung $|M| < |\mathcal{P}(M)|$ für den Spezialfall $M = V' = \{x \mid x \text{ ist Menge}\}$ und der Funktion $f = \text{id}_{V'}$, d. h. $f(x) = x$ für alle Mengen x . Aus dem Beweis des Satzes von Cantor wissen wir, daß $R = \{x \in V' \mid x \notin f(x)\}$ nicht im Wertebereich von f liegen kann (vgl. das Korollar zum Satz von Cantor in Kapitel 10). Aber $\text{rng}(f) = V'$, es gilt also $R \notin V'$. R ist aber nach dem Komprehensionsaxiom eine Menge, also haben wir $R \in V'$, Widerspruch!

Die Russell-Zermelo-Paradoxie ist also eine Reduzierung der Cantorschen Paradoxie auf einen Spezialfall. Dennoch ist sie von einer bestechenden Brillanz, und zudem ein rein logisches Paradoxon: Es wird ohne Zuhilfenahme von weitergehenden Sätzen gezeigt, daß ein R mit der Eigenschaft (+) nicht existieren kann, oder genauer, daß wir R nicht zu dem Bereich von Objekten rechnen dürfen, die wir in (+) für y einsetzen können. Das Cantorsche Paradoxon benutzt dagegen den nichttrivialen Satz von Cantor, und sein mengentheoretischer Inhalt ist, daß die Objekt- oder Mengenwelt selber keine Menge mehr ist – in dieser Formulierung sicher keine große Überraschung.

Eine einprägsame Version der Russell-Antinomie ist der „fleißige Barbier“: In einem Dorf lebt ein Barbier, der folgende Aussage macht: „Ich schneide genau denjenigen Dorfbewohnern die Haare, die sich ihre Haare nicht selbst schneiden.“ Nun ist der Barbier aber selber ein Dorfbewohner. Er muß sich also nach seiner Aussage die Haare genau dann schneiden, wenn er sie sich nicht selbst schneidet. Widerspruch! Der Barbier kann

also sein Versprechen nicht in die Tat umsetzen, seine Aussage ist eine Lüge. Außermathematische Beispiele für Objekte, die sich selbst als Element enthalten sind zudem etwa: Die „Menge aller Ideen“ ist wieder eine Idee; ein Katalog, der alle Titel von Büchern listet, listet seinen eigenen Titel; usw. In der heute üblichen axiomatischen Mengenlehre sind Mengen x mit der Eigenschaft „ $x \in x$ “ durch das sog. Fundierungsaxiom ausgeschlossen.

Übung

Sei $R^+ = \{x \mid x \notin x\} = R \cup \{x \mid x \text{ ist Grundobjekt}\}$.

Zeigen Sie die Paradoxie: $R^+ \in R^+ \text{ gdw } R^+ \notin R^+$.

Eine Variation der Russell-Zermelo-Paradoxie führt zur Paradoxie von Dimitry Mirimanov (1861 – 1945) aus dem Jahre 1917. Hierzu betrachten wir zunächst eine Verallgemeinerung der reflexiven Bedingung „ $x \in x$ “.

Definition (*n-reflexiv*)

Sei $n \in \mathbb{N}$. Eine Menge x heißt *n-reflexiv*, falls es Mengen $y_1, \dots, y_n$ gibt mit $x \in y_1 \in y_2 \in \dots \in y_n \in x$.

Insbesondere gilt: x ist 0-reflexiv gdw $x \in x$.

Übung

Für $n \in \mathbb{N}$ sei $R_n = \{x \mid x \text{ ist nicht } n\text{-reflexiv}\}$.

Zeigen Sie für alle $n \in \mathbb{N}$ die Paradoxie:

R_n ist n -reflexiv gdw R_n ist nicht n -reflexiv.

Statt $\in$ -Ketten, die bei x starten und wieder bei x enden, können wir auch unendlich absteigende $\in$ -Ketten betrachten. Mengen, für die solche Ketten nicht existieren, nennen wir fundiert:

Definition (*fundiert*)

Eine Menge x heißt *fundiert*, falls es keine Mengen $x_n, n \in \mathbb{N}$, gibt mit $x \ni x_0 \ni x_1 \ni x_2 \ni x_3 \dots$.

Ist $x \in x$, so ist x nicht fundiert, denn dann gilt $x \ni x \ni x \ni \dots$. Allgemeiner ist für $n \in \mathbb{N}$ jedes n -reflexive x nicht fundiert. Ist $x \ni x_0 \ni x_1 \ni \dots \ni x_n \ni \dots$, so ist offenbar nicht nur x , sondern auch jedes x_n selbst nicht fundiert.

Mirimanovsches Paradoxon, 1. Fassung

Sei $M = \{x \mid x \text{ ist fundiert}\}$.

Annahme, M ist fundiert.

Dann gilt $M \in M$, also ist M nicht fundiert. *Widerspruch*.

Also ist M nicht fundiert. Dann existieren Mengen $x_n, n \in \mathbb{N}$, mit

$M \ni x_0 \ni x_1 \ni x_2 \ni x_3 \ni \dots$

Dann ist aber x_0 nicht fundiert.

– Aber $x_0 \in M$, also ist x_0 fundiert nach Definition von M , *Widerspruch!*

Die Paradoxie läßt sich in einer Weise notieren, die an das Cantorsche Paradoxon erinnert. Dort ist $\mathcal{P}(V) \subseteq V$ die Schlüsseleigenschaft. Allgemein definieren wir:

Definition (*induktiv*)

Sei x eine Menge. x heißt *induktiv*, falls für alle y gilt:

Ist $y \subseteq x$, so ist $y \in x$.

Anders formuliert: x ist induktiv *gdw* $\mathcal{P}(x) \subseteq x$.

Wegen $x \subseteq x$ gilt $x \in x$ für alle induktiven x . Ist x induktiv und $y \subseteq x$, so ist auch $\mathcal{P}(y) \subseteq \mathcal{P}(x) \subseteq x$. Iteration ergibt, daß $\emptyset \subseteq \mathcal{P}(\emptyset) \subseteq \mathcal{P}(\mathcal{P}(\emptyset)) \subseteq \dots \subseteq x$, also sind $\emptyset, \mathcal{P}(\emptyset), \mathcal{P}(\mathcal{P}(\emptyset)), \dots$ Teilmengen und damit Elemente von allen induktiven Mengen x . Der Leser wird schnell sehen, daß induktive Mengen uferlos groß sind. Sie sind zudem stabil unter Durchschnitten, und dies führt zur zweiten Fassung der Mirimanov-Paradoxie.

Mirimanovsches Paradoxon, 2. Fassung

Sei $U = \bigcap \{ X \mid X \text{ ist induktiv} \}$.

Dann ist U induktiv.

Denn ist $x \subseteq U$, so ist $x \subseteq X$ für jedes induktive X .

Dann ist aber $x \in X$ für jedes induktive X , also $x \in U$.

Wegen U induktiv gilt $U \in U$.

Wir setzen

$U' = U - \{ U \}$.

Dann ist $U' \subset U$. Weiter ist U' induktiv.

Denn sei $x \subseteq U'$. Wegen $U' \subseteq U$ gilt dann $x \subseteq U$.

Also $x \in U$ wegen U induktiv.

Aber es gilt $x \neq U$ denn wir haben $U \notin x$, $U \in U$.

Also ist $x \in U - \{ U \} = U'$.

– Also $U' \subset U = \bigcap \{ X \mid X \text{ ist induktiv} \} \subseteq U'$, *Widerspruch!*

Die Sprechweise „erste und zweite Fassung“ ist gerechtfertigt, denn die in den Mirimanovparadoxien auftretenden „Mengen“ M und U sind identisch (und mit einem natürlichen Argument als identisch zu erkennen, das nicht direkt die paradoxalen Eigenschaften von M und U verwendet):

Übung

Sei $M = \{ x \mid x \text{ ist fundiert} \}$, $U = \bigcap \{ x \mid x \text{ ist induktiv} \}$.

Dann gilt $M = U$.

[zu $U \subseteq M$: Ist jedes $z \in y$ fundiert, so ist y fundiert. Also ist M induktiv.

zu $M \subseteq U$: Wir haben $U \subseteq M$. *Annahme*, es gibt ein $x_0 \in M - U$. Dann gilt

$\text{non}(x_0 \subseteq U)$ wegen U induktiv. Also existiert ein $x_1 \in x_0$ mit $x_1 \notin U$.

Induktiv zeigt man: Für alle n existiert ein $x_{n+1} \in x_n$ mit $x_{n+1} \notin U$.

Also ist $x_0 \in M$ nicht fundiert, *Widerspruch*.]

Interpretation der Paradoxien

Wir können die Paradoxien positiv deuten als:

Die Zusammenfassungen

$V = \{x \mid x = x\}$, $V' = \{x \mid x \text{ ist Menge}\}$, $R = \{x \in V' \mid x \notin x\}$,

$M = U = \{x \mid x \text{ ist fundiert}\} = \bigcap \{x \mid x \text{ ist induktiv}\}$

können wir nicht als Mengen, d. h. nicht als ein Ganzes, betrachten.

Bei der Cantorparadoxie entstand der Widerspruch durch die Anwendung des für alle Mengen gültigen Satzes $|M| < |\mathcal{P}(M)|$ auf $M = V$. Wir haben also gezeigt, daß der Satz nicht für V gilt, daß also V keine Menge ist. Ebenso ist V' keine Menge.

Im Falle der Paradoxie von Russell-Zermelo haben wir gezeigt, daß R nicht im Wertebereich V' der Identität $\text{id}_{V'}$ auf $V' = \{x \mid x \text{ ist Menge}\}$ liegt, d. h. R ist keine Menge.

Die Mirimanovsche Paradoxie zeigt, daß nicht nur der 0-irreflexive Teil R von V , sondern bereits der fundierte Teil M von V bzw. der Schnitt $U \subseteq V'$ aller induktiven Mengen nicht zu einem konsistenten Ganzen zusammengefaßt werden kann.

In diesen drei Fällen – und in allen weiteren aufgetretenen Paradoxien der Mengenlehre – ist für die Widersprüche lediglich das Komprehensionsaxiom verantwortlich, mit dessen Hilfe wir Zusammenfassungen unbegrenzt vieler Mengen zu einem Ganzen vornehmen dürfen. Dies führt zu Widersprüchen. Es gibt Grenzen der Methode, Vielheiten als Einheiten *oder* Mengen aufzufassen. Die Cantorsche Definition, die wir zu Beginn diskutiert hatten, erlaubt uns auch solche beliebige Zusammenfassungen gar nicht. Sie sagt: Eine Menge ist jede Vielheit, die als Ganzes aufgefaßt werden kann. Sie sagt nicht: Jede Vielheit ist eine Menge. Manche Vielheiten sind zu umfangreich, um iterativ als Objekte verwendet werden zu können. Wir können sie benennen, sie hinschreiben und sie uns vorstellen, aber sie zerfallen logisch, wenn wir sie zu Objekten machen wollen, und in die Mengenwelt, aus der sie von außen gezogen sind, zurückschicken. Der logische Zerfall des Russell-Zermelo-Konstrukts R als Objekt der Mengenwelt ist dabei instantan, er hängt nicht von mengentheoretischen Operationen wie der recht wilden Schnittbildung bei der Definition von U ab. R als Objekt der Mengenwelt nimmt sich selbst als Objekt auf und stößt sich selbst als Element „logisch gleichzeitig“ ab, ohne daß Potenzmengen, Vereinigungen oder ähnliches bei der Destruktion von R erst mithelfen müßten.

Cantor hat nur vereinzelt diskutiert, welche Vielheiten zu Mengen zusammengefaßt werden können (vgl. die Briefauszüge am Ende von 3.1). Er scheint Existenzannahmen und Bildungsprinzipien von Mengen als offenes Konzept verstanden zu haben. Seiner berühmten Ansicht, daß das Wesen der Mathematik gerade in ihrer Freiheit liege [1883 b, § 8], würde auch ein starres System von Mengenbildungsprinzipien nicht entsprechen. Wir untersuchen die Welt der Mengen, entdecken dabei immer neue Aspekte und lernen ihre Phänomene

immer besser kennen. Bei dieser Untersuchung des Mengenbegriffs stoßen wir auf die Ideen der Unendlichkeit, der Abzählbarkeit und Überabzählbarkeit, der Vergleichbarkeit, auf Eigenarten der reellen Zahlen und der Potenzmengenbildung, usw. Wir stoßen schließlich auch auf die Grenzen der „Zusammenfassung zu einem Ganzen,“ und dies ist eine weitere schöne Erkenntnis über die Mengenwelt, nicht etwa ein Hinweis auf ihren pathologischen Charakter. Die Idee der freien Entdeckung einer unabhängig von uns gegebenen Wirklichkeit bleibt durch die Paradoxien unberührt. Auch bei einem Erdbeben wird man nicht gleich an der Existenz des festen Bodens unter den Füßen zweifeln. Diese Sicht der Dinge stimmt schließlich auch mit der Geschichte der Mengenlehre überein: In ihrem Verlauf sind immer wieder neue Prinzipien über die Existenz von Mengen aufgestellt worden, deren genauere Erkundung mit „Entdeckung von Neu-land“ vielleicht am besten beschrieben wird.

Im Verlauf dieser Einführung haben wir viele Mengen gebildet, zum Teil recht umfassend im Übergang von M zu $\mathcal{P}(M)$, oder etwa im Beweis des Vergleichbarkeitssatzes. Es ist nun im Hinblick auf die Paradoxien der vollen Komprehension nur natürlich, sich – gewissermaßen als vorläufige Bestandsaufnahme – eine Liste von denjenigen Mengenbildungen zusammenzustellen, die wir für unsere Argumente wirklich brauchen, und die durch unsere Intuition und die bisherigen Erfahrungen mit dem Mengenbegriff als abgesichert, legitimiert und wünschenswert gelten. Daß zum Beispiel die Zusammenfassung aller Teilmengen einer unendlichen Menge zu einem Ganzen, der Potenzmenge, möglich ist, können wir nicht beweisen. Diese Zusammenfassung führt immerhin zu sehr großen fertigen Gesamtheiten, und es ist nicht auszuschließen, daß diese Operation widerspruchsvoll ist. Z. B. ist ein Widerspruch der folgenden Art denkbar: Wir bilden mit Hilfe von $\mathcal{P}(\mathbb{N})$ „diagonal“ eine Teilmenge von $\mathbb{N}$, die nicht in $\mathcal{P}(\mathbb{N})$ vorkommt. Demnach wäre die Zusammenfassung aller Teilmengen einer unendlichen Menge zu einem Ganzen nicht erlaubt, $\mathcal{P}(\mathbb{N})$ wäre, wie V oder R , keine Menge mehr, sondern nur eine Vielheit, die nicht als Ganzes betrachtet werden kann.

Bislang ist aber in der Liste der Mengenbildungen, die wir in dieser Einführung benutzt haben, und die wir im dritten Abschnitt explizit vorstellen werden, keine Paradoxie einer zu umfangreichen Zusammenfassung festgestellt worden. Die Liste schwebt aber prinzipiell immer in der Gefahr, substantiell verkleinert oder umgeschrieben werden zu müssen. Andererseits ist sie auch erweiterungsfähig, wenn sich unsere Kenntnis der Mengenwelt soweit entwickelt hat, daß wir eine eindeutige Erweiterung oder Erweiterungen in verschiedene Äste als natürlich und angemessen empfinden, – so natürlich und angemessen wie die Existenz von $\mathbb{N}$, $\mathbb{R}$, ω_1 , oder wie die Existenz der Potenzmenge einer beliebigen Menge.

In anderer Hinsicht sind abweichende Deutungen der Paradoxien denkbar. Die Interpretation, daß manche Komprehensionen „zu groß“ sind, drängt sich auf, und führt, viel wichtiger, zu einer natürlichen Theorie, die fast ganz wie die naive Mengenlehre ist. Vielleicht werden dieser Interpretation aber einmal andere, ebenso brauchbare und überzeugende an die Seite treten, die feinere Unterscheidungen bei der Komprehension treffen als nur „klein“ oder „groß“. Bis jetzt steht sie übermächtig und unerreicht hinter der Cantorsche Mengenlehre, und schützt sie mit der Keule vor allen irgendwie verdächtigen Instanzen der Komprehension.

Mengen und echte Klassen

V , V' und R sind, so können wir die Paradoxien deuten, zu groß, um Mengen zu sein. Im Fall von R heißt das positiv: Für viele Objekte gilt $x \notin x$.

Vielheiten, die keine Mengen mehr sein können, nennt man auch *echte Klassen*. Wir können dann als Ergebnis unserer Interpretation festhalten:

„Es gibt Mengen und es gibt echte Klassen.“

V , V' , R sind Beispiele für echte Klassen. Ob man echte Klassen als Objekte eines anderen Typs ansieht oder sie nur als Sprechweise betrachtet, also als eine Abkürzung für eine Vielheit der Form „alle x , die $\mathcal{E}(x)$ erfüllen“, und damit z. B. $x \in R$ als eine Abkürzung für „ x ist Menge und $x \notin x$ “, ist Geschmackssache. Die zweite Möglichkeit hat den Vorteil, daß wir neben den Grundobjekten und den Mengen nicht noch die Kategorie der echten Klassen einführen müssen. Wir werden diese Möglichkeit verfolgen, und echten Klassen keinen Objektstatus zukommen lassen. Da wir im dritten Abschnitt auch auf Grundobjekte verzichten, reden wir dann nur noch über Mengen, und „für alle x “ heißt dann immer „für alle Mengen x “, und „es gibt ein x “ heißt „es gibt eine Menge x “.

Der folgende Auszug aus einem Brief von Cantor an Dedekind zeigt, daß Cantor sich des Phänomens von Vielheiten, die keine Mengen mehr sind, voll bewußt war.

Cantor über konsistente und inkonsistente Vielheiten (Brief an Dedekind vom 3. August 1899):

„Hochverehrter Freund.

Wie ich Ihnen vor einer Woche schrieb, liegt mir viel daran, Ihr Urteil in gewissen fundamentalen Punkten der Mengenlehre zu erfahren und ich bitte Sie, die Ihnen dadurch verursachte Mühe mir zu verzeihen.

Gehen wir von dem Begriff einer bestimmten Vielheit (eines Systems, eines Inbegriffs) von Dingen aus, so hat sich mir die Notwendigkeit herausgestellt, zweierlei Vielheiten (ich meine immer *bestimmte* Vielheiten) zu unterscheiden.

Eine Vielheit kann nämlich so beschaffen sein, daß die Annahme eines ‚Zusammenseins‘ aller ihrer Elemente auf einen Widerspruch führt, so daß es unmöglich ist, die Vielheit als eine Einheit, als ‚ein fertiges Ding‘ aufzufassen. Solche Vielheiten nenne ich *absolut unendliche* oder *inkonsistente Vielheiten*.

Wie man sich leicht überzeugt, ist z. B. der ‚Inbegriff alles Denkbaren‘ eine solche Vielheit; später werden sich noch andere Beispiele darbieten.

Wenn hingegen die Gesamtheit der Elemente einer Vielheit ohne Widerspruch als ‚zusammenseiend‘ gedacht werden kann, so daß ihr Zusammengefaßtwerden zu ‚einem Ding‘ möglich ist, nenne ich sie eine *konsistente Vielheit* oder eine ‚Menge‘. (Im Französischen und Italienischen wird dieser Begriff durch die Worte ‚ensemble‘ und ‚insieme‘ treffend zum Ausdruck gebracht) ...“

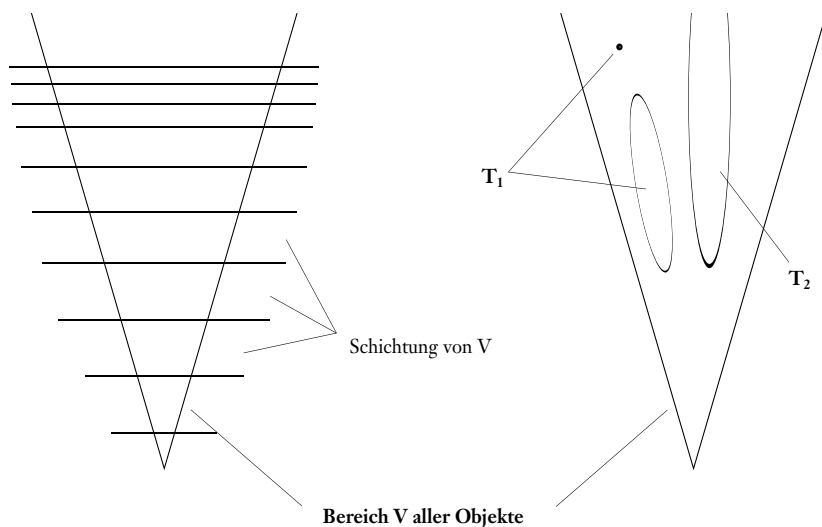
Cantor unterscheidet also zwischen Mengen und inkonsistenten Vielheiten; letztere als Mengen zu behandeln, sie etwa als Element einer Menge anzunehmen, führt zu Widersprüchen – und diese Widersprüche haben Cantor nicht beunruhigt, so wie sie uns heute nicht mehr beunruhigen.

Strukturierung des Bereichs aller Objekte durch Ränge

Eine anschauliche Vorstellung des Unterschiedes zwischen Mengen und echten Klassen kann man wie folgt erhalten. Man betrachtet den Bereich V aller Objekte. In geeigneter Weise wird nun jedem Objekt in V ein *Rang* zugeordnet. Ein Rang ist ein Maß für die Komplexität des Objekts, und je größer der Rang ist, desto komplizierter ist das Objekt; typischerweise ist dann z. B. $\mathbb{N}$ komplizierter als $\emptyset$, $\mathbb{R}$ komplizierter als $\mathbb{N}$, $\{\mathbb{R}, \mathbb{N}\}$ komplizierter als $\mathbb{R}$ und $\mathbb{N}$, $\mathcal{P}(M)$ immer komplizierter als M , usw. Man erhält so eine Schichtung von V . Echte Klassen sind dann genau die Teilbereiche von V , die Objekte von beliebig großem Rang enthalten, also mit unbeschränkt vielen Schichten nichtleeren Schnitt haben. Die Schichten selber sind demnach also immer Mengen.

Wie genau ein solcher Rang definiert werden kann, ist nicht trivial – es erfordert wieder die Ordinalzahlen. Die heute übliche Rangdefinition geht auf Mirimanov, John von Neumann (1903 – 1957) und Zermelo zurück [Mirimanov 1917a, von Neumann 1923, 1925, Zermelo 1930a]. (Wir geben diese Definition in 2.7.) Die im vorherigen Kapitel angegebene Schichtung von V in Bereiche gleicher Mächtigkeit ist als Rang nicht geeignet; denn jede Schicht – außer der Schicht, die nur die leere Menge enthält – ist eine echte Klasse. Dies kann man sich so klarmachen: Für jede Menge M ist $\{M\}$ ein Element der Schicht $\mathcal{S}_1$ aller Mengen mit genau einem Element. $\{M\}$ ist aber sicher mindestens so komplex wie M . $\mathcal{S}_1$ hat also Objekte mit beliebig großem Rang als Elemente, ist also eine echte Klasse.

Wir nehmen an, wir hätten ein geeignetes Maß für die Komplexität eines Objekts. Dann kann man folgendes Bild zeichnen:



Das Bild links zeigt die Schichtung von V gemäß einem geeigneten Rang, wobei kompliziertere Mengen in höheren Schichten liegen als einfachere. Die Frage ist nun, welche Teilbereiche von V Mengen sind.

Das Diagramm rechts zeigt einerseits einen im Rang beschränkten Teilbereich T_1 von V , der als Element einer höheren Schicht zugleich ein Element des Universums V ist. T_1 ist eine Menge. T_1 ist komplexer als jedes seiner Elemente, also liegt es in einer Schicht oberhalb von T_1 , aufgefaßt als Teilbereich des Universums V .

Zum anderen zeigt das Diagramm einen unbeschränkten Teilbereich T_2 von V , der in V nicht als Element vorkommen kann – T_2 ist eine echte Klasse: T_2 müßte als Menge in einer Schicht liegen, die oberhalb aller Schichten liegt, die T_2 – aufgefaßt als Teilbereich von V – trifft. Oberhalb des Teilbereichs T_2 von V gibt es aber im Gegensatz zu T_1 keine Schichten.

V selbst ist bei dieser Sichtweise trivialerweise eine echte Klasse, da V aus allen Schichten besteht und damit im Rang unbeschränkt ist.

Semantische Paradoxien

Schon einige Jahre vor dem Russell-Zermeloschen Paradoxon hatten Cantor und Burali-Forti mengentheoretische Paradoxien entdeckt, die mit den Ordinal- und Kardinalzahlen zusammenhängen; wir besprechen diese Paradoxien im zweiten Abschnitt. Durch die Veröffentlichung der Russell-Zermelo-Paradoxie durch Bertrand Russell 1903 wurden Paradoxien dann zum heißen Eisen der Logik. Selbstbezüglichkeit und Diagonalisierung hießen die Werkzeuge der paradoxen Ingenieurskunst, und ab 1905 wurde eine Reihe vorwiegend semantischer Paradoxien produziert. Die bekanntesten unter ihnen sind die Paradoxie von Jules Richard (1862 – 1956) aus dem Jahre 1905 [vgl. auch König 1905 b], die Paradoxie von „Mr. G. G. Berry of the Bodleian Library“, veröffentlicht mit dieser Fußnote von Russell 1908, sowie die Paradoxie von Kurt Grelling (1886 – 1942) in [Grelling / Nelson 1908]. Weit vor ihnen liegt die Paradoxie des Epimenides aus der Antike. Wir wollen diese Paradoxien kurz besprechen.

Paradoxie des Epimenides

Wir betrachten den Satz „Ich lüge jetzt.“ (oder „Diese Aussage ist falsch.“) Ist dieser Satz wahr oder falsch? Ist er wahr, so ist er falsch, und ist er falsch, so ist er wahr. (Die Aussage eines Kreters, der sagt „Alle Kreter lügen immer“ ist dagegen nicht paradox, sondern einfach eine Lüge.)

Oder: Sprechen Sie auf ein leeres Tonband mit zwei Stunden Laufzeit genau einen Satz bei Minute 60, und zwar diesen: „Alle Sätze auf diesem Tonband sind falsch.“ Spulen Sie das Band zurück, und hören Sie es sich ganz an. Haben Sie einen wahren Satz auf dem Tonband gehört oder nicht? (Das Band heißt „paradoxical meditation 120“.) Sprechen Sie den Satz nun 120 Mal in 120 Sprachen an verschiedenen Stellen auf ein anderes Band, und hören Sie sich das Band an. Haben Sie einen wahren Satz auf dem Tonband gehört?

(Das Band heißt „mankind searching for truth“.) Der Leser kann diese Gedankenexperimente weiter ausbauen, etwa mit vor- und zurückverweisenden Sätzen, wie zum Beispiel: „Der folgende Satz ist falsch.“ gefolgt von „Der vorhergehende Satz ist wahr.“ Daneben bilden Endlostonbänder eine gute Grundlage für logische Verwicklungen.

Die Paradoxie von Richard

Wir betrachten irgendeine (abzählbar unendliche) Liste $\mathfrak{S} = S_0, S_1, S_2, S_3, \dots, S_n, \dots$, $n \in \mathbb{N}$, aller Sätze der deutschen Sprache, die eine reelle Zahl definieren. Für $n \in \mathbb{N}$ sei $f(n)$ die reelle Zahl, die durch S_n definiert wird. Es sei dann $x = 0, b_0 b_1 b_2 \dots, b_n \in \{1, 2\}$ für $n \in \mathbb{N}$, die Cantorsche Diagonalisierung von $f(0), f(1), f(2), \dots$, (vgl. Kapitel 8). Dann gilt $x \neq f(n)$ für alle $n \in \mathbb{N}$. Aber es gilt, daß x definiert wird durch den Satz $S =$ „die Cantorsche Diagonalisierung der durch die Liste $\mathfrak{S}$ gegebenen reellen Zahlen $f(n)$, $n \in \mathbb{N}$.“ Also ist $S = S_n$ für ein n , und dann ist $x = f(n)$, *Widerspruch!*

Die Paradoxie von Berry

Sei $A \subseteq \mathbb{N}$ die Menge der natürlichen Zahlen, die durch einen Satz der deutschen Sprache definiert werden können, der höchstens 19 Wörter lang ist. Dann ist A endlich, da es nur endlich viele derartige Sätze gibt. Sei also $n = \min(\mathbb{N} - A)$. Dann gilt $n =$ „die kleinste natürliche Zahl, die nicht durch einen Satz der deutschen Sprache mit höchstens neunzehn Wörtern definiert werden kann“. Dann ist n aber durch einen Satz mit 19 Wörtern definierbar, also gilt $n \in A$, *Widerspruch!*

Die Paradoxie von Grelling

Das Wort „blau“ ist nicht blau, aber das Wort „mehrsilbig“ ist mehrsilbig, „deutsch“ ist deutsch, „kalt“ ist nicht kalt, aber „abstrakt“ ist abstrakt. Wir nennen ein Wort *selbsteinschließend* (oder *prädikabel*), falls das Wort unter den durch das Wort bezeichneten Begriff fällt, und *selbstausschließend* sonst. Die Frage ist nun: Ist „selbstausschließend“ selbsteinschließend oder selbstausschließend? Ist „selbstausschließend“ selbsteinschließend, so trifft es auf sich selbst zu und ist also selbstausschließend. Ist aber „selbstausschließend“ selbstausschließend, so trifft es auf sich zu, also ist es selbsteinschließend, *Widerspruch!*

Die Paradoxien sind in dieser Form zwar unterhaltsam, aber nicht wirklich bedrohlich oder mathematisch ernst zu nehmen, denn sie reden nicht über klar definierte mathematische Objekte. Und wenn man die Paradoxien außermathematisch betrachtet, gewinnt man den Eindruck, daß sie verschiedene Sprachebenen vermischen, oder Sprache und Definierbarkeit als etwas Absolutes betrachten. Dem Auge des Lesers wird geraten, nicht zu lange dem unaufhörlichen „wahr-falsch“-Pendel der Paradoxien zu folgen. Logische Spitzfindigkeiten zehren, wie bei Ovid die durchwachten Nächte, an den Kräften junger Männer (und Frauen).

Interessanterweise lassen sich aber die Ideen hinter diesen Paradoxien durch Formalisierung mathematisch umsetzen, und sie führen dann nicht zu Widersprüchen, sondern zu fundamentalen Ergebnissen: Der erste Gödelsche Unvollständigkeitssatz gipfelt in der Konstruktion einer formalen Aussage, die inhaltlich besagt: „ich bin nicht beweisbar“ oder „diese Aussage ist nicht beweisbar“. Ein Satz von Tarski lautet: Es gibt keine arithmetische Definition der Wahrheit von arithmetischen Sätzen (und das Gleiche gilt innerhalb der Mengenlehre für die Wahrheit von Sätzen der Mengenlehre; dagegen gibt es in der Mengenlehre eine mengentheoretische Definition der arithmetischen Wahrheit). In der Berechenbarkeitstheorie zeigt man: Es gibt keine effektive Aufzählung aller berechenbaren Funktionen $f: \mathbb{N} \rightarrow \mathbb{N}$. Die Unentscheidbarkeit des Halteproblems gehört auch hierher: Es gibt kein Programm, das andere Programme einliest, und dann entscheidet, ob das eingelesene Programm, wenn man es anwirft, terminiert oder nicht. Weiter kann man beweisen: Es gibt eine effektiv auflistbare Menge $A \subseteq \mathbb{N}$, für welche die Frage „ist $n \in A$?“ nicht algorithmisch beantwortbar ist. Die Menge aller (Kodes von) beweisbaren Sätzen der axiomatischen Zahlentheorie – oder einer axiomatischen Mengenlehre – ist ein Beispiel für eine solche listbare, aber nicht entscheidbare Menge. Hinter all diesen mathematischen Sätzen steht ein Diagonalargument. Es ist klar, daß man schon für die bloße mathematische Formulierung solcher Resultate sehr sorgfältig vorgehen muß, und wir müssen den Leser hier auf die Lehrbücher zur mathematischen Logik verweisen.

Abraham Fraenkel über Georg Cantor

„Was die *Persönlichkeit* C.s im allgemeinen betrifft, so berichten alle, die ihn kannten, von seinem sprühenden, witzigen, originellen Naturell, das leicht zur Explosion neigte und stets voll heller Freude über die eigenen Einfälle war; von dem niemals ermüdenden Temperament, das die Teilnahme seiner auch äußerlich imponierenden, großen Gestalt an einer Mathematikerversammlung zu einem ihrer lockendsten Reize machte, das bis in die späte Nacht wie auch in früher Morgenstunde seine Gedanken (zu seinen mathematischen und den vielseitigen außermathematischen Interessengebieten) förmlich überquellen ließ; von seinem lauterem Charakter, treu seinen Freunden, hilfreich, wo es nötig war, liebenswürdig im Verkehr; nebenbei auch von einer typischen Gelehrtenzerstreutheit. Im mündlichen wissenschaftlichen Gedankenaustausch war er mehr der Gebende; es lag ihm nicht, unmittelbar vorgetragene fremde Ideen sogleich aufzufassen. All seinen Gedanken war er mit der gleichen Liebe und Intensität hingegeben; in stärkerem Maße vielleicht noch als der aufgewandte Scharfsinn und selbst als die mit begrifflicher Gestaltungskraft gepaarte geniale Intuition ist die ungeheure Energie, mit der er seine Gedanken über alle Hindernisse und Hemmungen hinweg verfolgte und an ihnen festhielt, das Instrument gewesen, dem wir die Entstehung der Mengenlehre zu danken haben. Solch unerschütterliche Zähigkeit entsprang seiner tiefen Überzeugung von der Wahrheit, ja Wirklichkeit seiner Ideen...

Einen der großen Bahnbrecher der Wissenschaft hat die mathematische Welt, und zu unserem Stolz speziell auch unsere Deutsche Mathematikervereinigung, in Georg Cantor besessen. Die allgemeine Verbreitung der Erkenntnis, daß sein Werk der Analysis

1. Transfinite Operationen

Wir geben in den ersten Kapiteln dieses Abschnitts eine Einführung in die Theorie der Ordinalzahlen. Es handelt sich hierbei um die Fortsetzung der natürlichen Zahlenreihe

$$0, 1, 2, 3, \dots$$

ins Transfinite:

$$(+)\ 0, 1, 2, \dots, \omega, \omega + 1, \omega + 2, \dots, \omega + \omega, \omega + \omega + 1, \dots, \dots, \alpha, \alpha + 1, \dots, \dots, \dots$$

Die Elemente dieser Reihe heißen *Ordinalzahlen*, oder, wie Cantor sie genannt hat, *Ordnungszahlen*. Die natürlichen Zahlen bilden ein Anfangsstück der Ordinalzahlen. Die erste Ordinalzahl, die größer ist als alle natürlichen Zahlen, wird seit Cantor mit ω bezeichnet, ein Zeichen, das an das Unendlichkeitssymbol erinnert. Ab einschließlich ω werden die Elemente der Reihe *transfinite Zahlen* genannt.

Die Idee der Ordinalzahlen ist die folgende: Von irgendeiner Stelle der Reihe können wir eine beliebige Reise nach rechts antreten, Schritt für Schritt, oder auch, indem wir große Abschnitte überspringen. Jede solche Reise hat eine eindeutig bestimmte Ordinalzahl als Ziel, den Limes oder das Supremum aller Schritte. Von dieser Ordinalzahl können wir nun eine neue Reise starten – die Ordinalzahlen sind in diesem Sinne unerschöpflich. Die Länge einer solchen Wanderung entlang der Ordinalzahlen ist dabei auch nicht auf die natürlichen Zahlen beschränkt, sondern wird konsequenterweise wieder durch eine Ordinalzahl beschrieben.

Anhand der reellen Zahlengeraden kann man sich diese Idee zumindest ein erstes Stück weit vor Augen führen. Wir betrachten hierzu etwa die Menge $M = \mathbb{N} \cup \{n - 1/k \mid n, k \in \mathbb{N}, n \geq 1, k \geq 2\}$. M enthält alle natürlichen Zahlen, und zu jeder natürlichen Zahl ungleich Null eine unendliche Folge, die diese Zahl von links approximiert. Versucht man nun, die Elemente von M von links nach rechts durchzuzählen, so wird man zwangsläufig die Schritte dieser Abzählung mit einer Reihe wie in (+) bezeichnen:

Schritt	0	1	2	...	ω	$\omega + 1$	$\omega + 2$	...	$\omega + \omega$	...
Element von M	0	$1 - 1/2$	$1 - 1/3$	...	1	$2 - 1/2$	$2 - 1/3$	...	2	...

Die Menge M visualisiert alle endlichen und transfiniten Ordinalzahlen $\omega \cdot n + m$ für $n, m \in \mathbb{N}$, wobei $\omega \cdot 0 = 0$, $\omega \cdot 1 = \omega$, $\omega \cdot 2 = \omega + \omega$, ..., $\omega \cdot (n + 1) = \omega \cdot n + \omega$. Der Limes aller zur Aufzählung von M benötigten Schritte hat die natürliche Be-

zeichnung $\omega \cdot \omega$, er entspricht aber keinem Element der Menge M mehr. Komplizierter ist die folgende Teilmenge der reellen Zahlen:

Übung

Sei $M' = \mathbb{N} \cup \{n - 1/k - 1/m \mid n, k, m \in \mathbb{N}, n \geq 1, k \geq 2, m \geq k(k-1)\}$.

Zählen Sie M' in einer Tabelle auf, analog zur obigen für M .

Weisen Sie den dazu benötigten Schritten geeignete ω -Symbole zu, etwa $\omega \cdot \omega$, $(\omega \cdot \omega) + \omega$, ..., usw.

Wir betrachten hier ω zunächst nur als ein Symbol, dessen Bedeutung es ist, einen Schritt zu bezeichnen, der sich an unendlich viele, mit den natürlichen Zahlen bezeichnete Schritte anschließt. Eine präzise Definition von ω und den anderen transfiniten Zahlen geben wir im weiteren Verlauf dieses Abschnitts, nachdem wir uns den Begriff der Ordinalzahl inhaltlich vertraut gemacht haben.

Wir werden zeigen, daß dem Zählprozeß über die natürlichen Zahlen hinaus nichts Vages oder Mystisches anhaftet, und daß die Ordinalzahlen mit gutem Recht als Zahlen bezeichnet werden können. Mit ihrer Hilfe können wir die Elemente jeder beliebigen endlichen oder unendlichen Menge durchzählen und numerieren. Wir zählen die Elemente einer Menge entlang obiger Reihe (+) auf, bis die Menge M erschöpft ist.

Es zeigt sich, daß die Reihe der Ordinalzahlen unermesslich lang ist. Wie „weit draußen“ eine Ordinalzahl α liegen kann, ist bis heute Gegenstand der mengentheoretischen Forschung (wobei hier nicht das Fehlen einer unstrittigen Definition der Ordinalzahlen gemeint ist, sondern Fragen der Form „Existiert eine Ordinalzahl mit den und den Eigenschaften?“). In jedem Falle gibt es aber ungeheuer viele Ordinalzahlen: Wie wir sehen werden, bilden die Ordinalzahlen eine echte Klasse – es gibt ihrer zu viele, als daß sie ein „fertiges Ganzes“ bilden könnten. In einer geeigneten Schichtung des Universums V durch Ränge kann man sie sich als – im Rang unbeschränkte – senkrechte Mittelachse von V vorstellen.

Eine formale Begründung der Ordinalzahltheorie, des „Gehirns der Mengenlehre“, innerhalb einer axiomatischen Festung ist sicher wünschenswert, zumal man durch die Ordinalzahlen in natürlicher Weise an die Grenzen des mengentheoretischen Universums gelangt, wo die Physik eine andere wird. Die Ideen und Intuitionen hinter den Ordinalzahlen sind aber, wie im Fall der Mächtigkeitstheorie, unabhängig von einem formalen Rahmen, und gehen ihm unbedingt voraus.

In Übereinstimmung mit der historischen Entwicklung gesellt sich einer informellen Diskussion der Ordinalzahlen in natürlicher Weise ein weiteres Themengebiet der Cantorsche Forschung hinzu, nämlich die Untersuchung von Teilmengen reeller Zahlen – „linearen Punktmannigfaltigkeiten“ in den Worten Cantors. Seine Arbeiten in diesem Gebiet wurden zum Ausgangspunkt zweier weiterer neuer mathematischer Disziplinen des 20. Jahrhunderts, der Topologie und der Maßtheorie, und die Begriffe, die wir hier anhand der reellen Zahlen entwickeln, stehen heute in sehr allgemeiner Form in den Eingangshallen des

mathematischen Gesamtgebäudes. Innerhalb der Mengenlehre werden die Cantorschen Ideen zur Untersuchung der reellen Zahlen von der *deskriptiven Mengenlehre* fortgeführt.

Drei Ansatzpunkte für transfinite Zahlen

Das Zählen über die natürlichen Zahlen hinaus hat sich bisher bereits an mehreren Stellen fast von selbst aufgedrängt:

- (1) Beim Algorithmus des Abtragens zum Vergleich der Größe zweier Mengen. Wenn dieser Algorithmus für unendliche Mengen durchgeführt wird, werden nach dem Abtragen von unendlich vielen Elementen unter Umständen Zwischenstufen, sogenannte Limeschritte, nötig (vgl. 1.4 und die Diskussion nach dem Beweis des Vergleichbarkeitssatzes in 1.5). Wir brauchen also Stufen des Abtragens hinter den natürlichen Zahlen, und diese Stufen sind offenbar von der Form obiger Reihe (+).
- (2) Bei der Erzeugung immer größerer Mächtigkeiten durch iterierte Anwendung der Potenzmengenoperation (vgl. 1.10 und 1.11):

$$\mathcal{P}(M), \mathcal{P}^2(M), \mathcal{P}^3(M), \dots, M' = \bigcup_{n \in \mathbb{N}} \mathcal{P}^n(M), \mathcal{P}(M'), \mathcal{P}^2(M'), \dots, \dots$$

Auch hier entspricht die Struktur der durchgeführten Operationen der Reihe (+), wobei wir „im Nachfolgerschritt“ die Potenzmengenoperation anwenden, und „im Limeschritt“ die Vereinigung der bisherigen Reihe bilden.
- (3) Bei der Schichtung von V durch Ränge (vgl. 1.13). Die ersten Schichten nummerieren wir der Reihe nach mit den natürlichen Zahlen 0, 1, 2, 3, ... Alle diese Schichten sind Mengen. Die Vereinigung abzählbar vieler Mengen ist aber sicher wieder eine Menge, und somit gibt es eine Schicht nach den ersten unendlich vielen Schichten, und auch nach dieser Schicht kommen weitere Schichten. Die Nummern der Schichten entsprechen wieder der Reihe (+).

Der Leser, der die zweite Hälfte des Kapitels über Kardinalzahlarithmetik nicht zurückgestellt hat, kennt ein weiteres Beispiel: Kardinalzahlen haben einen Nachfolger, und Mengen von Kardinalzahlen haben ein Supremum. Die Auflistung der Kardinalzahlen nach ihrer Größe beginnt mit 0, 1, 2, ..., $\aleph_0$, $(\aleph_0)^+$, ... Sie hat die Struktur der Reihe (+): Einer Kardinalzahl α folgt α^+ unmittelbar nach, und an eine Menge $\mathfrak{A}$ von Kardinalzahlen ohne ein größtes Element schmiegt sich „im Limeschritt“ $\mathfrak{b} = \sup(\mathfrak{A})$ an.

Bevor wir aus den Gemeinsamkeiten dieser drei Punkte die Ordinalzahlen herauslösen, behandeln wir im nächsten Kapitel ein weiteres Beispiel – dasjenige Beispiel, an Hand dessen Cantor die Ordinalzahlen entdeckte. Einer der ersten Erfolge Cantors als Mathematiker war der Beweis eines Eindeutigkeitssatzes über die Entwicklung einer Funktion in eine trigonometrische Reihe. Erweiterungen dieses Satzes – im Hinblick auf zulässige Ausnahmemengen für die Kon-

vergenz der Reihe – führten Cantor zur Untersuchung bestimmter Teilmengen reeller Zahlen. Diese Untersuchung nahm bald einen abstrakten Charakter an und entwickelte eine Eigendynamik, aus der schließlich neben der Mengenlehre auch die Topologie und die Maßtheorie hervorgehen sollten.

Wir beenden dieses Kapitel mit einer Bemerkung von Felix Hausdorff, die eine Haltung ausdrückt, die unserer Unbekümmertheit gegenüber den aufgetretenen Paradoxien entspricht: Diese Paradoxien bedürfen zwar langfristig einer Klärung, jedoch ist dies im Hinblick auf eine inhaltliche Weiterentwicklung der Cantorschen Mengenlehre nicht vorrangig.

Felix Hausdorff über die Überbewertung der Paradoxien

„Daß eine Untersuchung wie diese, die den positiven Bestand der noch so jungen Mengenlehre im Sinne ihres Schöpfers um einen, wenn auch nur bescheidenen, Zuwachs zu vermehren trachtet, sich nicht prae limine damit aufhalten kann, in die Diskussion um die Prinzipien der Mengenlehre einzutreten, wird vielleicht an den Stellen Anstoß erregen, wo gegenwärtig ein etwas deplaziertes Maß von Scharfsinn an diese Diskussion verschwendet wird. Einem Beobachter, der es auch der Skepsis gegenüber nicht an Skepsis fehlen läßt, dürften die ‚finitistischen‘ Einwände gegen die Mengenlehre ungefähr in drei Kategorien zerfallen: in solche, die das ernsthafte Bedürfnis nach einer, etwa axiomatischen Verschärfung des Mengenbegriffs verraten; in diejenigen, die mitsamt der Mengenlehre die ganze Mathematik treffen würden, endlich in einfache Absurditäten einer an Worte und Buchstaben sich klammernden Scholastik. Mit der ersten Gruppe wird man sich heute oder morgen verständigen können, die zweite darf man getrost auf sich beruhen lassen, die dritte verdient schärfste und unzweideutigste Ablehnung. In der vorliegenden Arbeit werden diese drei Reaktionen stillschweigend vollzogen ...“

(Felix Hausdorff 1908, „Grundzüge einer Theorie der geordneten Mengen“)

2. Lineare Punktmengen

In diesem Kapitel ist $\mathbb{R}$ die zugrunde liegende Struktur, und wir betrachten beliebige Teilmengen P von $\mathbb{R}$, die wir wie Cantor auch als „lineare Punktmengen“ bezeichnen.

Wir halten vorab zwei wesentliche Eigenschaften der reellen Zahlen fest, die uns schon begegnet sind.

(1) $\mathbb{R}$ ist vollständig.

$P \subseteq \mathbb{R}$ heißt *nach oben beschränkt*, falls ein $s \in \mathbb{R}$ existiert mit $P \leq s$, d. h. es gilt $x \leq s$ für alle $x \in P$. Ein solches s heißt eine *obere Schranke* von P .

Analog heißt P *nach unten beschränkt*, falls ein $s \in \mathbb{R}$ existiert mit $s \leq P$, d. h. es gilt $s \leq x$ für alle $x \in P$. Ein solches s heißt eine *untere Schranke* von P .

P heißt *beschränkt*, falls P nach oben und unten beschränkt ist.

Die *Vollständigkeit* von $\mathbb{R}$ lautet nun:

Jede nichtleere nach oben beschränkte Teilmenge P von $\mathbb{R}$ besitzt ein *Supremum* (kleinste obere Schranke), d. h. es gibt ein $s^* \in \mathbb{R}$ mit:

(1) $P \leq s^*$

(2) Ist $s \in \mathbb{R}$ und $P \leq s$, so gilt $s^* \leq s$.

s^* wird mit $\sup(P)$ bezeichnet und heißt das Supremum von P ; es ist eindeutig bestimmt.

Es folgt, daß jede nichtleere nach unten beschränkte Teilmenge P von $\mathbb{R}$ ein *Infimum* (größte untere Schranke) besitzt, d. h. es gibt ein $s^* \in \mathbb{R}$ mit:

(1) $s^* \leq P$

(2) Ist $s \in \mathbb{R}$ und $s \leq P$, so gilt $s \leq s^*$.

s^* wird mit $\inf(P)$ bezeichnet und heißt das Infimum von P ; es ist eindeutig bestimmt.

[Zum Beweis der Existenz von Infima setze man $s^* = \sup(\{s \in \mathbb{R} \mid s \leq P\})$.]

Suprema und Infima beschränkter Teilmengen P von $\mathbb{R}$ können Elemente von P sein oder nicht. So ist z. B.

$\sup(\{x \in \mathbb{R} \mid x \leq 1\}) = \sup(\{x \in \mathbb{R} \mid x < 1\}) = 1$.

(2) $\mathbb{R}$ hat eine abzählbare dichte Teilmenge: $\mathbb{Q}$ ist dicht in $\mathbb{R}$.

Hierzu eine allgemeine Definition:

Definition (*dichte Teilmenge*)

Sei $P \subseteq \mathbb{R}$. P heißt *dicht in* $\mathbb{R}$, falls gilt:

Für alle $a, b \in \mathbb{R}$ mit $a < b$ existiert ein $x \in P$ mit $a < x < b$.

Der Leser, der die Übung „ $\mathbb{Q}$ ist dicht in $\mathbb{R}$ “ in Abschnitt 1, Kapitel 7 ausgelassen hat, möge sie jetzt nachholen.

Aus „ $\mathbb{Q}$ ist dicht in $\mathbb{R}$ “ folgt, daß jede reelle Zahl als Supremum (oder Infimum) einer Menge von rationalen Zahlen darstellbar ist. Man kann umgekehrt diese Eigenschaft verwenden, um die reellen Zahlen aus den rationalen Zahlen zu konstruieren.

Häufungspunkte

Cantor hat zum ersten Mal über die natürlichen Zahlen hinausgezählt, als er *Ableitungen* von Teilmengen von $\mathbb{R}$ studierte. Diese Ableitungen streifen alle verlorenen Schafe von Punktmengen ab. Für die mathematische Definition brauchen wir eine Reihe von einfachen Begriffen.

Definition (reelle Intervalle)

Ein $I \subseteq \mathbb{R}$ heißt ein (*reelles*) *Intervall*, falls für alle $a, b \in I$ gilt:

Ist $c \in \mathbb{R}$ mit $a < c < b$, so ist $c \in I$.

Ein nichtleeres Intervall I heißt:

- (i) *nach links (rechts) geschlossen*, falls $\inf(I) \in I$ ($\sup(I) \in I$),
- (ii) *nach links (rechts) offen*, falls $\inf(I)$ ($\sup(I)$) nicht existiert oder $\inf(I) \notin I$ ($\sup(I) \notin I$),
- (iii) *offen (geschlossen)*, falls I nach links und rechts offen (geschlossen) ist.

Das Intervall $I = \emptyset$ betrachten wir als ein Intervall jedes Typs.

Für $a, b \in \mathbb{R}$ setzen wir:

$$]a, b[= \{x \in \mathbb{R} \mid a < x < b\}, \quad [a, b] = \{x \in \mathbb{R} \mid a \leq x \leq b\},$$

$$[a, b[= \{x \in \mathbb{R} \mid a \leq x < b\}, \quad]a, b] = \{x \in \mathbb{R} \mid a < x \leq b\}.$$

Jedes beschränkte Intervall können wir in dieser Form schreiben. Für unbeschränkte Intervalle verwenden wir die Darstellungen $] -\infty, a]$, $] -\infty, a[$, $[a, +\infty[$, $[a, +\infty]$, $] -\infty, +\infty[$. Für $b \leq a$ ist $]a, b[= \emptyset$, und für alle $a \in \mathbb{R}$ ist $[a, a] = \{a\}$.

Beschränkte offene Intervalle sind darüber hinaus durch ihren Mittelpunkt und ihre Ausdehnung festgelegt. Diese Form wird häufig gebraucht, und es ist nützlich, einen Begriff für sie zur Verfügung zu stellen.

Definition (ε -Umgebung)

Sei $\varepsilon \in \mathbb{R}$, $\varepsilon > 0$. Für $a \in \mathbb{R}$ setzen wir:

$$U_\varepsilon(a) =]a - \varepsilon, a + \varepsilon[.$$

$U_\varepsilon(a)$ heißt die (*offene*) ε -*Umgebung* [epsilon-Umgebung] von a .

Übung

Seien $a, b, \varepsilon_1, \varepsilon_2 \in \mathbb{R}$, $\varepsilon_1, \varepsilon_2 > 0$. Sei $U = U_{\varepsilon_1}(a) \cap U_{\varepsilon_2}(b)$. Dann ist U leer oder eine ε -Umgebung eines Punktes.

Damit können wir nun Häufungspunkte von Teilmengen von $\mathbb{R}$ definieren:

Definition (*Häufungspunkt und isolierter Punkt*)

Sei $P \subseteq \mathbb{R}$.

- (i) $a \in \mathbb{R}$ heißt *Häufungspunkt von P*, falls jede ε -Umgebung von a Punkte aus $P - \{a\}$ enthält, d.h. falls gilt:
Für alle $\varepsilon > 0$ ist $U_\varepsilon(a) \cap (P - \{a\}) \neq \emptyset$.
- (ii) $a \in P$ heißt *isolierter Punkt von P*,
falls a kein Häufungspunkt von P ist.

Häufungspunkte von P müssen also nicht Elemente von P sein, während „ a ist isolierter Punkt von P “ impliziert, daß $a \in P$ gilt.

Cantor hat in seiner Arbeit „Über die Ausdehnung eines Satzes aus der Theorie der trigonometrischen Reihen“ von 1872 die folgende Definition gegeben:

Cantor (1872): „Unter einem Grenzpunkt [Häufungspunkt] einer Punktmenge P verstehe ich einen Punkt der Geraden von solcher Lage, daß in jeder Umgebung desselben unendlich viele Punkte aus P sich befinden, wobei es vorkommen kann, daß er außerdem selbst zur Menge gehört. Unter Umgebung eines Punktes sei aber hier ein jedes Intervall verstanden, welches den Punkt *in seinem Inneren* hat.“

Übung

Zeigen Sie die Äquivalenz der Cantorschen Definition eines Grenzpunktes von P mit obiger Definition eines Häufungspunktes von P .

Isolierte Punkte einer Menge können nicht durch andere Elemente der Menge beliebig gut approximiert werden. Wir können einen isolierten Punkt in ein Intervall einschließen, das disjunkt vom Rest der Menge ist.

Seien $P = \{1/n \mid n \in \mathbb{N}, n > 0\}$, $Q = P \cup \{0\}$. Für jedes $n \in \mathbb{N}$, $n \neq 0$, ist dann $1/n$ ein isolierter Punkt von P und Q , und 0 ist jeweils der einzige Häufungspunkt der beiden Mengen.

Übung

Sei $P \subseteq \mathbb{R}$ beschränkt und nichtleer. Zeigen Sie:

Ist $\sup(P) \notin P$, so ist $\sup(P)$ ein Häufungspunkt von P .

Analog ist $\inf(P)$ ein Häufungspunkt von P , falls $\inf(P) \notin P$.

Ein wichtiger Existenzsatz über Häufungspunkte ist der Satz von Bolzano-Weierstraß.

Satz (*Satz von Bolzano-Weierstraß*)

Sei P eine unendliche beschränkte Teilmenge von $\mathbb{R}$.

Dann existiert ein Häufungspunkt von P .

Beweis

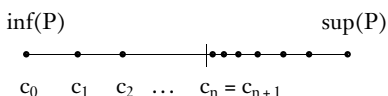
Wir definieren rekursiv $c_0 \leq c_1 \leq c_2 \leq \dots$ durch:

$$c_0 = \inf(P),$$

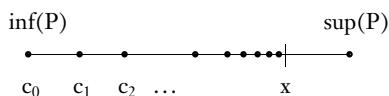
$$c_{n+1} = \inf(P - [c_0, c_n]) \text{ für } n \in \mathbb{N}.$$

Wegen P unendlich ist c_n definiert für alle $n \in \mathbb{N}$.

Ist $c_n = c_{n+1}$ für ein n , so ist c_n ein Häufungspunkt von P (!), vgl. Übung oben.



1. Fall



2. Fall

Andernfalls ist $c_0 < c_1 < c_2 < \dots$ und $C = \{c_n \mid n \in \mathbb{N}\} \subseteq P$ (!).

Wegen $C \subseteq P$ und P beschränkt existiert $x = \sup(C)$,

– und x ist ein Häufungspunkt von P .

Einige Beispiele zur im Beweis konstruierten Folge $c_0, c_1, \dots$ sind: Ist $P = [0, 1]$ oder $P =]0, 1]$, so ist $c_n = 0$ für alle $n \in \mathbb{N}$, und 0 ist Häufungspunkt von P . Ist $P = \{1 - 1/n \mid n \in \mathbb{N}, n \geq 1\}$, so ist $c_n = n/(n+1)$ für alle $n \in \mathbb{N}$, und der konstruierte Häufungspunkt von P ist 1.

Cantor (1872): „... Darnach ist es leicht zu beweisen, daß eine aus einer unendlichen Anzahl von Punkten bestehende [beschränkte] Punktmenge stets zum Wenigsten einen Grenzpunkt hat.“

Grenzwerte von Folgen

Als Korollar zum Satz von Bolzano-Weierstraß zeigen wir noch, daß jede Folge reeller Zahlen, in der die Abstände beliebig weit auseinanderliegender Folgenglieder beliebig klein werden, einen eindeutigen Häufungspunkt oder *Limes* besitzt. Hierzu einige Begriffe.

Definition (*konvergente Folgen und Limes einer Folge*)

Eine Folge $x_0, x_1, x_2, \dots$ reeller Zahlen heißt *konvergent gegen* $x \in \mathbb{R}$, in Zeichen $\lim_{n \rightarrow \infty} x_n = x$, falls gilt:

Für alle $\varepsilon > 0$, $\varepsilon \in \mathbb{R}$, existiert ein $n_0 \in \mathbb{N}$ mit: $x_n \in U_\varepsilon(x)$ für alle $n \geq n_0$.

x heißt dann der *Limes* oder *Grenzwert* der Folge $x_0, x_1, \dots$

Wir begnügen uns häufig mit Schreibweisen der Form $x_0, x_1, x_2, \dots$ für Folgen. Andere Schreibweisen sind $(x_n)_{n \in \mathbb{N}}$ oder in der Mengenlehre oft auch das noble $\langle x_n \mid n \in \mathbb{N} \rangle$. Offiziell ist eine *Folge in einer Menge* M eine Funktion $f: \mathbb{N} \rightarrow M$, und x_n ist dann nur eine andere Schreibweise für $f(n)$.

Übung

Der Limes einer konvergenten Folge $x_0, x_1, x_2, \dots$ ist eindeutig bestimmt, d.h. für alle $x, y \in \mathbb{R}$ gilt:

$$\lim_{n \rightarrow \infty} x_n = x \text{ und } \lim_{n \rightarrow \infty} x_n = y \text{ folgt } x = y.$$

Konvergente Folgen verdichten sich also an genau einer Stelle. Dies hat zur Konsequenz, daß die Glieder der Folge beliebig eng zusammenrücken. Diese Eigenschaft wird präzise gefaßt im Begriff der Cauchyfolge.

Definition (*Cauchyfolge*)

Eine Folge $x_0, x_1, x_2, \dots$ reeller Zahlen heißt *Cauchyfolge* oder *Fundamentalfolge* in $\mathbb{R}$, falls gilt:

Für alle $\varepsilon > 0$, $\varepsilon \in \mathbb{R}$ existiert ein $n_0 \in \mathbb{N}$ mit:

$$|x_n - x_m| < \varepsilon \text{ für alle } n, m \in \mathbb{N} \text{ mit } n, m \geq n_0.$$

Jede konvergente Folge ist eine Cauchyfolge: Ist $\varepsilon > 0$ und $|x - x_n| < \varepsilon/2$ für alle $n \geq n_0$, so ist $|x_n - x_m| < \varepsilon$ für alle $n, m \geq n_0$. Aber auch die Umkehrung ist richtig, und der wesentliche Teil des Beweises dieser Behauptung wird durch den Satz von Bolzano-Weierstraß getragen:

Satz (*Konvergenz von Cauchyfolgen*)

Jede Cauchyfolge $x_0, x_1, x_2, \dots$ konvergiert.

Beweis

Sei $X = \{x_n \mid n \in \mathbb{N}\}$.

1. *Fall*: X ist endlich.

Sei dann $X = \{y_1, \dots, y_k\}$ mit $y_1 < y_2 < \dots < y_k$, $k \geq 1$.

Ist $k = 1$, so ist offenbar y_1 der Limes der Folge. Sei also $k > 1$.

Sei $\varepsilon = \min \{|y_i - y_{i+1}| \mid 1 \leq i < k\}$.

Weiter sei $n_0 \in \mathbb{N}$ mit $|x_n - x_m| < \varepsilon$ für alle $n, m \geq n_0$.

Nach Wahl von ε existiert dann ein $i \leq k$ mit $x_n = y_i$ für alle $n \geq n_0$.

Also gilt $\lim_{n \rightarrow \infty} x_n = y_i$.

2. *Fall*: X ist unendlich.

X ist beschränkt, denn es gibt ein n_0 mit $|x_n - x_{n_0}| < 1$ für alle bis auf endlich viele $n \in \mathbb{N}$. Nach dem Satz von Bolzano-Weierstraß existiert also ein Häufungspunkt x von X . Wir zeigen $\lim_{n \rightarrow \infty} x_n = x$.

Sei hierzu $\varepsilon > 0$, $\varepsilon \in \mathbb{R}$, und sei $n_0 \in \mathbb{N}$ mit

$$|x_n - x_m| < \varepsilon/2 \text{ für alle } n, m \geq n_0.$$

Wegen x Häufungspunkt von X ist $X \cap U_{\varepsilon/2}(x)$ unendlich.

Also existiert ein $m \geq n_0$ mit $|x - x_m| < \varepsilon/2$.

Für alle $n \geq n_0$ ist dann

$$|x - x_n| = |x - x_m + x_m - x_n| \leq |x - x_m| + |x_m - x_n| < \varepsilon/2 + \varepsilon/2 = \varepsilon.$$

— Also ist $x_n \in U_\varepsilon(x)$ für alle $n \geq n_0$.

Ableitungen

Cantor hat die Häufungspunkte einer Teilmenge von $\mathbb{R}$ als Kennzeichen für die Reichhaltigkeit der Menge betrachtet, und hierzu den Begriff der Ableitung einer Punktmenge eingeführt.

Definition (*Cantorsche Ableitung einer Menge reeller Zahlen*)

Sei $P \subseteq \mathbb{R}$. Wir setzen

$$P' = \{x \in \mathbb{R} \mid x \text{ ist Häufungspunkt von } P\}.$$

P' heißt die (*Cantorsche*) *Ableitung von P*.

P' entsteht aus P durch Säumen des Randes von P und anschließender Entfernung der isolierten Punkte.

Cantor (1872): „Es ist nun ein bestimmtes Verhalten eines jeden Punktes der Geraden zu einer gegebenen Menge P , entweder ein Grenzpunkt derselben oder kein solcher zu sein, und es ist daher mit der Punktmenge P die Menge ihrer Grenzpunkte begrifflich mit gegeben, welche ich mit P' bezeichnen und *die erste abgeleitete Punktmenge* von P nennen will.“

Elementare Eigenschaften der Ableitung sind:

Übung

Seien $P, Q \subseteq \mathbb{R}$. Dann gilt:

(i) $P \subseteq Q$ folgt $P' \subseteq Q'$.

(ii) $(P \cup Q)' = P' \cup Q'$.

(iii) $(P \cap Q)' \subseteq P' \cap Q'$.

[Dagegen ist $P' \cap Q' \subseteq (P \cap Q)'$ im allgemeinen nicht richtig.]

(iv) Ist P endlich, so ist $P' = \emptyset$.

Wir betrachten einige Beispiele:

(i) $\emptyset' = \emptyset, \mathbb{R}' = \mathbb{R},$

(ii) $]0, 1[' = [0, 1], \mathbb{Q}' = \mathbb{R},$

(iii) $(\{1/n \mid n \in \mathbb{N}, n > 0\} \cup \{0\})' = \{0\},$

(iv) $\{1/n \mid n \in \mathbb{N}, n > 0\}' = \{0\}.$

Es sind also die Fälle

(i) $P' = P, \quad$ (ii) $P \subset P', \quad$ (iii) $P' \subset P, \quad$ (iv) $P' \cap P = \emptyset$

möglich. Das Beispiel $\mathbb{Q}' = \mathbb{R}$ zeigt, daß die Ableitung einer Menge von größerer Mächtigkeit sein kann als die Ausgangsmenge.

Übung

Sei $P \subseteq \mathbb{R}$. Dann sind äquivalent:

- (i) P ist dicht in $\mathbb{R}$. (ii) $P' = \mathbb{R}$.

Abgeschlossene, in sich dichte und perfekte Mengen

Das Verhältnis von P' zu P führt zu drei natürlichen Begriffen:

Definition (*abgeschlossene und perfekte Teilmengen von $\mathbb{R}$*)

Ein $P \subseteq \mathbb{R}$ heißt:

- (i) *abgeschlossen*, falls $P' \subseteq P$,
- (ii) *in sich dicht*, falls $P \subseteq P'$,
- (iii) *perfekt*, falls $P' = P$.

„Abgeschlossen“ heißt also für eine Menge: Jeder Häufungspunkt der Menge gehört bereits zur Menge.

Cantor (1884b): „Wenn eine Punktmenge P so beschaffen ist, daß ihre Ableitung $P^{(1)}$ [= P'] in ihr als Divisor [Teilmenge] enthalten ist, oder was dasselbe ist, daß

$$\mathfrak{D}(P, P^{(1)}) [= P \cap P^{(1)}] = P^{(1)},$$

so wollen wir P eine *abgeschlossene Menge* nennen.“

„In sich dicht“ bedeutet: Jeder Punkt der Menge läßt sich beliebig gut durch andere Punkte der Menge approximieren, oder anders: es gibt keine isolierten Punkte.

Cantor (1884b): „Es ist ferner wichtig, den Fall ins Auge zu fassen, daß eine Menge P Divisor ihrer Ableitung $P^{(1)}$ ist oder, was dasselbe ist, daß

$$\mathfrak{D}(P, P^{(1)}) [= P \cap P^{(1)}] = P;$$

unter solchen Umständen wollen wir P eine *in sich dichte Menge* nennen.“

Eine Menge ist perfekt, wenn sie abgeschlossen ist und keine isolierten Punkte besitzt. Die abgeschlossenen Intervalle sowie $\mathbb{R}$ sind Beispiele für perfekte Mengen. Kompliziertere perfekte Mengen, die keine Intervalle als Teilmengen enthalten, werden wir später diskutieren, wenn wir uns ausführlicher mit der Cantormenge beschäftigen, die uns im ersten Abschnitt schon begegnet ist (1.9).

Cantor (1883b): „... S dagegen ist so beschaffen, daß bei dieser Punktmenge der Ableitungsprozeß gar keine Änderung hervorbringt, indem

$$S = S^{(1)} [= S']$$

... ist; derartige Mengen S nenne ich *perfekte Punktmengen*.“

Der zur Abgeschlossenheit duale Begriff der offenen Menge wurde erst 1902 von Henri Lebesgue (1875 – 1941) eingeführt:

Definition (*offene Menge*)

Ein $Q \subseteq \mathbb{R}$ heißt *offen*, falls $\mathbb{R} - Q$ abgeschlossen ist.

Eine fundamentale Eigenschaft der offenen Mengen ist in der folgenden Charakterisierung gegeben:

Übung

Sei $Q \subseteq \mathbb{R}$. Dann sind äquivalent:

- (i) Q ist offen.
- (ii) Für alle $a \in Q$ existiert ein $\varepsilon > 0$ mit $U_\varepsilon(a) \subseteq Q$.



Eine offene Menge enthält also um jeden ihrer Punkte eine ε -Umgebung, und die Offenheit einer Menge folgt umgekehrt aus dieser Eigenschaft, die von „abgeschlossen“ und damit vom Begriff der Ableitung keinen Gebrauch mehr macht. Man kann sie zur Definition von „offen“ verwenden, und die abgeschlossenen Mengen dann als die Komplemente der offenen Mengen einführen. In dieser Weise werden die Begriffe in der Topologie der reellen Zahlen heute üblicherweise behandelt.

Nichtleere offene Mengen enthalten also immer nichtleere offene Intervalle. Wegen $]a, b[=]\mathbb{R}[$ für alle $a, b \in \mathbb{R}$ mit $a < b$ können wir damit die Mächtigkeit offener Mengen sofort angeben:

Satz (*Mächtigkeit offener Mengen*)

Sei $Q \subseteq \mathbb{R}$ offen und nichtleer.

Dann gilt $|Q| = |\mathbb{R}|$.

Die möglichen Mächtigkeiten der abgeschlossenen Mengen zu bestimmen ist wesentlich schwieriger. Jede endliche Teilmenge von $\mathbb{R}$ sowie $\mathbb{N}$ und $\mathbb{R}$ selbst sind abgeschlossene Teilmengen der reellen Zahlen. Alle endlichen Mächtigkeiten sowie „abzählbar unendlich“ und „gleichmächtig zu $\mathbb{R}$ “ sind also mögliche Größen von abgeschlossenen Mengen. Aus dem Satz von Cantor–Bendixson (Kapitel 11) wird folgen, daß die abgeschlossenen Mengen keine weiteren Mächtigkeiten besitzen. Die Kontinuumshypothese gilt also für die offenen und für die abgeschlossenen Mengen: Jede offene oder abgeschlossene Menge ist abzählbar oder gleichmächtig zu $\mathbb{R}$.

Die offenen Mengen sind stabil unter beliebigen Vereinigungen und endlichen Durchschnitten. Für abgeschlossene Mengen gelten die dualen Eigenschaften:

Übung

- (i) Sei $\mathcal{Q} \subseteq \mathcal{P}(\mathbb{R})$ und jedes $Q \in \mathcal{Q}$ sei offen.
Dann ist $\bigcup \mathcal{Q}$ offen. Ist $\mathcal{Q}$ endlich, so ist $\bigcap \mathcal{Q}$ offen.
- (ii) Sei $\mathcal{Q} \subseteq \mathcal{P}(\mathbb{R})$ und jedes $Q \in \mathcal{Q}$ sei abgeschlossen.
Dann ist $\bigcap \mathcal{Q}$ abgeschlossen. Ist $\mathcal{Q}$ endlich, so ist $\bigcup \mathcal{Q}$ abgeschlossen.

Wir untersuchen nun die Operation der Ableitung genauer. Wir haben gesehen, daß bei der Ableitung einer Punktmenge neue Punkte hinzukommen können. Ist dagegen die Punktmenge selbst die Ableitung einer Punktmenge, so kann die Ableitung die Punktmenge nur noch verkleinern. Anders: Die Ableitung einer Punktmenge ist immer abgeschlossen.

Satz (*Abgeschlossenheit der Ableitung*)

Sei $P \subseteq \mathbb{R}$. Dann ist P' abgeschlossen.

Beweis

Wir müssen $(P')' \subseteq P'$ zeigen.

Die Idee ist: Ein Häufungspunkt von Häufungspunkten einer Menge P ist selbst ein Häufungspunkt von P .

Sei also $a \in (P')'$ beliebig. Sei weiter $\varepsilon > 0$. Wir müssen zeigen:

$$(+)\ U_\varepsilon(a) \cap (P - \{a\}) \neq \emptyset.$$

Wegen $a \in (P')'$ existiert ein $b \in P'$ mit $b \neq a$, $b \in U_{\varepsilon/2}(a)$.

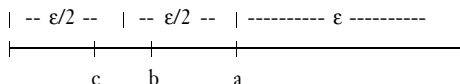
Sei $\delta = |b - a|$. Dann gilt

$$\delta > 0 \text{ und } U_\delta(b) \subseteq U_\varepsilon(a), a \notin U_\delta(b).$$

Wegen $b \in P'$ existiert dann ein $c \in P$ mit $c \in U_\delta(b)$.

Dann gilt $c \in P$, $c \in U_\varepsilon(a)$, $c \neq a$.

— Dies zeigt (+) und damit $a \in P'$ für alle $a \in (P')'$.



Cantor (1884b): „Jede Menge, welche selbst erste Ableitung einer anderen Menge ist, gehört auch, wie wir wissen, zu den abgeschlossenen Mengen.“

Definition (*Abschluß und Inneres einer Punktmenge*)

Sei $P \subseteq \mathbb{R}$. Wir setzen:

$$\text{cl}(P) = P \cup P', \quad [\text{cl für engl. closure}]$$

$$\text{int}(P) = \mathbb{R} - \text{cl}(\mathbb{R} - P). \quad [\text{int für engl. interior}]$$

$\text{cl}(P)$ heißt *der Abschluß von P*, $\text{int}(P)$ heißt *das Innere von P*.

Übung

- (i) Für alle $P \subseteq \mathbb{R}$ ist $\text{cl}(P)$ die kleinste abgeschlossene Obermenge von P :
 $\text{cl}(P)$ ist abgeschlossen, und ist $Q \supseteq P$ abgeschlossen, so ist $\text{cl}(P) \subseteq Q$.
- (ii) Für alle $P \subseteq \mathbb{R}$ ist $\text{int}(P)$ die größte offene Teilmenge von P .
Es gilt $\text{int}(P) = \bigcup \{ U_\varepsilon(x) \subseteq P \mid x \in P, \varepsilon > 0 \}$.

Wir haben gezeigt, daß P' immer abgeschlossen ist, d. h. es gilt $(P')' \subseteq P'$ für alle Punktmengen P . Die Ableitung einer Punktmenge ist aber im allgemeinen nicht perfekt:

Übung

- (i) Konstruieren Sie ein $P \subseteq \mathbb{R}$ mit $(P')' \subset P'$.
- (ii) Ist P in sich dicht, so ist P' perfekt.

Im allgemeinen haben wir also mit der Bildung der Ableitung P' von P keine lineare Punktmenge erreicht, die perfekt, d. h. stabil gegenüber der Bildung der Ableitung ist; $(P')'$ kann eine echte Teilmenge von P' sein. Es ist nun nur natürlich, die Folge $P, P', P'' = (P')', P''' = (P'')', \dots$ zu betrachten. Dies führt zu *iterierten Ableitungen* und damit letztendlich fast zwangsläufig zu den Ordinalzahlen.

Iterierte Ableitungen

Definition (iterierte Ableitung)

Sei $P \subseteq \mathbb{R}$. Dann ist $P^{(n)}$, die n -te Ableitung von P , für $n \in \mathbb{N}$ rekursiv definiert durch:

$$P^{(0)} = P,$$

$$P^{(n+1)} = (P^{(n)})' \quad \text{für } n \in \mathbb{N}.$$

Cantor (1879b): „Da hiernach die Ableitung einer Punktmenge P wieder eine bestimmte Punktmenge P' ist, so kann auch von dieser die Ableitung gesucht werden, welche alsdann *zweite Ableitung* von P genannt und mit P'' bezeichnet wird; durch eine Fortsetzung dieses Verfahrens erhält man die v -te Ableitung von P , welche mit $P^{(v)}$ bezeichnet wird.“

Nach dem Satz oben ist es nur für die erste Ableitung $P' = P^{(1)}$ möglich, keine Teilmenge der zugrunde liegenden Menge zu sein. Alle weiteren Ableitungen sind abgeschlossen und daher gilt

$$P^{(1)} \supseteq P^{(2)} \supseteq P^{(3)} \supseteq P^{(4)} \supseteq \dots$$

Cantor (1879b): „Bemerkenswert ist ferner, daß alle Punkte von $P'', P''', \dots$ auch immer Punkte von P' sind, während ein zu P' gehöriger Punkt nicht notwendig auch ein solcher von P ist.“

Eine natürliche Frage ist:

Terminiert diese Folge immer in einer perfekten Menge, d. h. gibt es für alle $P \subseteq \mathbb{R}$ immer ein $n \in \mathbb{N}$ mit der Eigenschaft $P^{(n)} = P^{(n+1)}$?

Die Antwort ist nein:

Übung

Konstruieren Sie ein abgeschlossenes $P \subseteq \mathbb{R}$ mit

$$P \supset P^{(1)} \supset P^{(2)} \supset P^{(3)} \supset \dots,$$

$$\text{d.h. } P^{(n+1)} \subset P^{(n)} \text{ für alle } n \in \mathbb{N}.$$

Eine unendliche $\supset$ -absteigende Kette von iterierten Ableitungen ist also möglich. Andererseits sind alle $P^{(n)}$ abgeschlossene Mengen, falls P abgeschlossen ist. Der Durchschnitt einer solchen absteigenden Kette ist wieder nichtleer und abgeschlossen, falls P beschränkt ist:

Übung

Seien P_n beschränkte, abgeschlossene und nichtleere Teilmengen von $\mathbb{R}$ mit $P_{n+1} \subseteq P_n$ für alle $n \in \mathbb{N}$.

Dann ist $\bigcap_{n \in \mathbb{N}} P_n$ nichtleer und abgeschlossen.

[Für „nichtleer“ wird die Vollständigkeit von $\mathbb{R}$ benutzt; es gilt $\sup \{ \inf(P_n) \mid n \in \mathbb{N} \} \in \bigcap_{n \in \mathbb{N}} P_n$.]

Abgeschlossene und beschränkte Teilmengen von $\mathbb{R}$ heißen auch *kompakt*.

Es kann also weitere nichttriviale Stufen im Ableitungsprozeß geben.

Definition

Sei $P \subseteq \mathbb{R}$. Dann ist $P^{(\omega)}$, die ω -te Ableitung von P , definiert durch

$$P^{(\omega)} = \bigcap_{n \in \mathbb{N}} P^{(n)}.$$

Cantor (1880d): „... so wird P' sich aus zwei wesentlich verschiedenen Punktmengen Q und R zusammensetzen [$P' = Q \cup R$, $Q \cap R = \emptyset$] ... Q besteht aus denjenigen Punkten von P' , die bei hinreichendem Fortschreiten in der Folge $P', P'', P''', \dots$ verloren gehen, die andere R umfaßt diejenigen Punkte, welche in *allen* Gliedern der Folge $P', P'', P''', \dots$ erhalten bleiben, es ist also R definiert durch die Formel:

$$R = \mathfrak{D}(P', P'', P''', \dots).$$

Wir haben aber auch offenbar:

$$R = \mathfrak{D}(P'', P''', P^{IV}, \dots)$$

und allgemein:

$$R = \mathfrak{D}(P^{(n_1)}, P^{(n_2)}, P^{(n_3)}, \dots),$$

wo $n_1, n_2, n_3, \dots$ irgend eine Reihe ins Unendliche wachsender ganzer positiver Zahlen ist.

Diese aus der Menge P hervorgehende Punktmenge R werde nun durch das Zeichen:

$$P^{(\infty)}$$

ausgedrückt und Ableitung von P der Ordnung ∞ genannt.“

Später hat Cantor das Zeichen ω für ∞ benutzt, das eher ein festes Objekt suggeriert, und doch zugleich an das Unendlichkeitssymbol erinnert.

Haben wir $P^{(\omega)}$ gebildet, so können wir weiter die Ableitungen $P^{(\omega)'} \supseteq P^{(\omega)''} \supseteq P^{(\omega)'''} \supseteq \dots$ betrachten. Man kann wieder Mengen $P \subseteq \mathbb{R}$ konstruieren, für welche

alle hierbei auftretenden Inklusionen echt sind. In diesem Fall haben wir also auch in $(\omega + \omega)$ -vielen Schritten noch keine perfekte Menge im Ableitungsprozeß erreicht, also immer noch keinen Index α gefunden für den $P^{(\alpha)} = (P^{(\alpha)})'$ gilt. Sobald wir ein perfektes $P^{(\alpha)}$ erreichen, wird der weitere Ableitungsprozeß trivial, alle folgenden Schritte sind dann identisch mit $P^{(\alpha)}$. Andernfalls fallen bei den weiteren Schritten immer wieder Punkte aus den iterierten Ableitungen heraus. Wir setzen:

$$\begin{aligned} P^{(\omega+n)} &= (P^{(\omega)})^{(n)} \quad \text{für } n \in \mathbb{N}, \text{ und} \\ P^{(\omega+\omega)} &= (P^{(\omega)})^{(\omega)} = \bigcap_{n \in \mathbb{N}} P^{(\omega+n)}. \end{aligned}$$

Damit erhalten wir eine Kette:

$$P^{(1)} \supseteq P^{(2)} \supseteq \dots \supseteq P^{(n)} \supseteq \dots \supseteq P^{(\omega)} \supseteq P^{(\omega+1)} \supseteq P^{(\omega+2)} \supseteq \dots \supseteq P^{(\omega+n)} \supseteq \dots \supseteq P^{(\omega+\omega)}.$$

Die Inklusionen können wieder echt sein – und es gibt dann weitere nichttriviale Schritte der Ableitung $P^{(\omega+\omega)'}$, $P^{(\omega+\omega)''}$, ...

Hier liegt also wieder eine Situation vor wie in (+) und den drei Beispielen des ersten Kapitels. Alle vier Beispiele verweisen auf die Idee der Ordinalzahlen und zeigen ihre Eigenart, die Stufen solcher transfiniter Prozesse markieren zu können, so wie die natürlichen Zahlen die Stufen gewöhnlicher unendlicher Folgen indizieren. Eine mathematische Umsetzung dieser Idee erscheint nun wünschenswert, denn wir hätten dann eine nicht nur schöne, sondern auch vielseitig einsetzbare Erweiterung der natürlichen Zahlen gewonnen. Transfinite Prozesse tauchen, wie wir gesehen haben, in vielerlei Situationen in natürlicher Weise auf.

Zur Einführung der Ordinalzahlen kann man sich ein Stück weit mit einer rein symbolischen Bezeichnung der Stufen behelfen, und dabei die vertrauten arithmetischen Operationen verwenden. Cantor hat bereits 1880 solche elementar arithmetischen Stufen angegeben. Statt von transfiniten Zahlen oder Ordnungszahlen spricht er damals noch von Unendlichkeitssymbolen.

Cantor (1880d): „Die erste Ableitung von $P^{(\infty)}$ werde mit $P^{(\infty+1)}$, die n^{te} Ableitung von $P^{(\infty)}$ mit $P^{(\infty+n)}$ bezeichnet; $P^{(\infty)}$ wird aber auch eine, im Allgemeinen von $O[\emptyset]$ verschiedene Ableitung von der Ordnung ∞ haben, wir nennen sie $P^{(2\infty)}$. Durch Fortsetzung dieser Begriffskonstruktionen kommt man zu Ableitungen, die konsequenterweise durch:

$$P^{(n_0 \infty + n_1)}$$

zu bezeichnen sind, wo n_0, n_1 positive ganze Zahlen sind. Wir kommen aber auch darüber hinaus, indem wir:

$$\mathfrak{D}(P^{(\infty)}, P^{(2\infty)}, P^{(3\infty)}, \dots)$$

bilden und dafür das Zeichen $P^{(\infty^2)}$ festsetzen.

Hieraus ergibt sich durch Wiederholung derselben Operation und Kombinierung mit den früher gewonnenen der allgemeinere Begriff:

$$P^{(n_0 \infty^2 + n_1 \infty + n_2)},$$

und durch Fortsetzung dieses Verfahrens kommt man zu:

$$P^{(n_0 \infty^v + n_1 \infty^{v-1} + \dots + n_v)},$$

wo $n_0, n_1, \dots, n_v$ positive ganze Zahlen sind. Zu weiteren Begriffen gelangt man, indem man v variabel werden läßt; man setze:

$$P^{(\infty)} = \mathfrak{D} (P^{(\infty)}, P^{(\infty^2)}, P^{(\infty^3)}, \dots).$$

Durch konsequentes Fortschreiten gewinnt man sukzessive die weiteren Begriffe:

$$P^{(n \infty^\infty)}, P^{(\infty \infty^{+1})}, P^{(\infty^\infty + n)}, P^{(\infty^{n \infty})}, P^{(\infty^\infty^\infty)}, P^{(\infty^\infty^\infty^\infty)} \text{ u. s. w.};$$

wir sehen hier eine dialektische Begriffserzeugung*), welche immer weiter führt und dabei frei jeglicher Willkür in sich notwendig und konsequent bleibt ...

[Fußnote] *): Zu derselben bin ich vor nun zehn Jahren gelangt; bei Gelegenheit einer eigentümlichen Darstellung des Zahlbegriffs (Math. Ann. Bd. V) habe ich entfernt darauf hingewiesen.“

Wir haben die Fußnote hier mitaufgenommen, weil sie zeigt, wie lange völlig neuartige Ideen brauchen können, um klar ins Bewußtsein zu kommen. In der erwähnten Arbeit (1872b) hatte Cantor die reellen Zahlen aus sogenannten Fundamentalfolgen, bestehend aus rationalen Zahlen konstruiert. Durch Iteration dieses Verfahrens gelangt er zu reellen Zahlen immer höherer (endlicher) Art, die dem Wert nach mit den üblichen reellen Zahlen übereinstimmen, von deren Konstruktionsprozeß er sich aber Nutzen für die Analysis erhoffte – eine Hoffnung, die sich nicht erfüllt hat. Ein wirklich transfiniten Prozeß findet sich in der Arbeit noch nicht; den Keim zu den Ordinalzahlen trägt sie aber sicherlich in sich, zumal Cantor auch hier schon die Ableitungen $P, P', P'', \dots$ betrachtet hat.

Wir können nun die Reihe (+) aus dem ersten Kapitel ergänzen durch die von Cantor eingeführten „Unendlichkeitssymbole“. Wir schreiben hierbei wie heute üblich ω für ∞ und weiter $\omega n = \omega \cdot n$ für $n \infty$.

Die Reihe lautet dann:

$$\begin{aligned} (+) \quad & 0, 1, 2, \dots, \omega, \omega + 1, \dots, \omega \cdot 2 = \omega + \omega, \omega \cdot 2 + 1, \dots, \omega \cdot 3, \dots, \omega \cdot 4, \dots, \\ & \omega \cdot \omega = \omega^2, \omega^2 + 1, \dots, \omega^2 + \omega, \dots, \omega^2 \cdot 2, \dots, \omega^3, \dots, \omega^\omega, \dots, \omega^{\omega^\omega}, \dots, \alpha, \dots \end{aligned}$$

Übung

Der Leser überlege sich die in (+) unterdrückten „Limesstellen“, z. B. $\omega^2 + \omega 2$.

Wir fragen wieder: Wie groß kann α werden? Wie weit kommt man nach rechts?

Die Antwort ist: Wir haben mit obigen durch arithmetische Operationen bezeichneten Ordinalzahlen erst einen winzigen Bruchteil der transfiniten Zahlen kennengelernt! Denn: Wir können immer weiter zählen, und jedem Weg nach rechts einen Limes hinzufügen, so wie wir den natürlichen Zahlen den Limes ω hinzugefügt haben. Solche Erweiterungen werden der Idee der Ordinalzahlen nur gerecht.

Man muß nun zudem, egal wie weit man die Reihe fortsetzt, schnell aufgeben, die Ordinalzahlen durch arithmetische Operationen oder anderswie konkret zu

bezeichnen – wir werden sehen, daß die Ordinalzahlen eine echte Klasse bilden; es gibt ihrer zu viele, um ihnen allen einen Namen geben zu können. Aber so wie man irgendwann mit dem Ausdruck „sei $n \in \mathbb{N}$ “ etwas anfangen kann, entwickelt man durch Abstraktion und Gewöhnung im Laufe der Zeit ein Gefühl für den Ausdruck „sei α eine Ordinalzahl“.

Wir brauchen noch eine mathematische Definition von „ α ist eine Ordinalzahl“, denn wir wollen die Ordinalzahlen nicht, wie etwa die natürlichen Zahlen, als Grundobjekte behandeln. Wir geben eine solche Definition im Stil von Cantor und Hausdorff in den nächsten Kapiteln, und kommen schließlich zur modernen, allen Maßstäben an Strenge genügenden Definition. Entscheidend ist der Begriff der Wohlordnung, den wir im nächsten Kapitel ausführlich behandeln. Erst mit diesem Begriff gelingt es, alle Elemente der Reihe (+) ein für allemal in einer Definition einzufangen.

Mit Hilfe der Ordinalzahlen können wir dann auch die hier offen gebliebenen Fragen beantworten:

- (1) Wie viele Ableitungen $P^{(\alpha)}$ braucht man, um von einer Menge P zu einer perfekten Menge zu gelangen, d. h. zu einem $P^{(\alpha)}$ mit $P^{(\alpha)} = P^{(\alpha+1)}$?
- (2) Welche Mächtigkeiten sind für die abgeschlossenen Mengen möglich?

Zu (1) bemerken wir wieder, daß die Zahlen aus obiger Liste nicht genügen, wovon man sich mit einiger Mühe überzeugen kann.

Frage (2) ist eine Approximation an das Kontinuumsproblem. Anstatt alle Teilmengen von $\mathbb{R}$ zu betrachten, untersucht man interessante Teilmengen von $\mathcal{P}(\mathbb{R})$, und versucht zu zeigen, daß alle ihre Elemente „regulären“, „nicht-pathologischen“ Charakter besitzen. Die einfachsten Beispiele für solche reguläre Teilmengen von $\mathcal{P}(\mathbb{R})$ sind die abgeschlossenen und die offenen Mengen. Dieses Programm ist von der deskriptiven Mengenlehre in der Nachfolge Cantors sehr erfolgreich weiterverfolgt worden, und die Ergebnisse dieser Teildisziplin der Mengenlehre bilden ein Gegengewicht zur Unabhängigkeit der Kontinuums-hypothese und anderer Aussagen über die reellen Zahlen.

Georg Cantor über die Einführung der transfiniten Zahlen

„Die bisherige Darstellung meiner Untersuchungen in der Mannigfaltigkeitslehre¹⁾ [Fußnote 1 enthält die in 1.1 bei der Diskussion des Platonismus wiedergegebene frühe Mengendefinition *Vieles, welches sich als Eines denken läßt.*] ist an einen Punkt gelangt, wo ihre Fortführung von einer Erweiterung des realen ganzen Zahlbegriffs $[\mathbb{N}]$ über die bisherigen Grenzen hinaus abhängig wird, und zwar fällt diese Erweiterung in eine Richtung, in welcher sie meines Wissens bisher von Niemandem gesucht worden ist.

Die Abhängigkeit, in welche ich mich von dieser Ausdehnung des Zahlbegriffs versetzt sehe, ist eine so große, daß es mir ohne letztere kaum möglich sein würde, zwanglos den kleinsten Schritt weiter vorwärts in der Mengenlehre auszuführen; möge in diesem Umstande eine Rechtfertigung oder, wenn nötig, eine Entschuldigung dafür gefunden werden, daß ich scheinbar fremdartige Ideen in meine Betrachtungen einführe. Denn es handelt sich um eine Erweiterung resp. Fortsetzung der realen ganzen Zahlenreihe über das

Unendliche hinaus; so gewagt dies auch scheinen möchte, kann ich dennoch nicht nur die Hoffnung, sondern die feste Überzeugung aussprechen, daß diese Erweiterung mit der Zeit als eine durchaus einfache, angemessene, natürliche wird angesehen werden müssen. Dabei verhehle ich mir keineswegs, daß ich mit diesem Unternehmen in einen gewissen Gegensatz zu weitverbreiteten Anschauungen über das mathematische Unendliche und zu häufig vertretenen Ansichten über das Wesen der Zahlgröße mich stelle.

Was das mathematische Unendliche anbetrifft, soweit es eine berechnete Verwendung in der Wissenschaft bisher gefunden und zum Nutzen derselben beigetragen hat, so scheint mir dasselbe in erster Linie in der Bedeutung einer veränderlichen, entweder über alle Grenzen hinaus wachsenden oder bis zu beliebiger Kleinheit abnehmenden, aber stets endlich bleibenden Größe aufzutreten. Ich nenne dieses Unendliche das Uneigentlich-Unendliche.

... Die unendlichen realen ganzen Zahlen, welche ich im Folgenden definieren will und zu denen ich schon vor einer längeren Reihe von Jahren geführt worden bin, ohne daß es mir zum deutlichen Bewußtsein gekommen war, in ihnen konkrete Zahlen von realer Bedeutung zu besitzen ..., haben durchaus nichts gemein mit dem ... Uneigentlich-Unendlichen, dagegen ist ihnen derselbe Charakter der Bestimmtheit eigen, wie wir ihn bei dem unendlich fernen Punkte in der analytischen Funktionentheorie antreffen; sie gehören also zu den Formen und Affektionen des Eigentlich-Unendlichen ... “

(Georg Cantor 1883b, „Über unendliche lineare Punktmannigfaltigkeiten (V)“)

3. Wohlordnungen

Unsere vier Beispiele führten uns zu einer Reihe der Form:

$$(+)\quad 0, 1, 2, 3, \dots, \omega, \omega + 1, \omega + 2, \omega + 3, \dots, \dots, \alpha, \dots, \dots, \dots$$

Wir konnten diese Reihe nur ein Stück weit „von unten“ beschreiben. Die Aufgabe ist nun, auf einen Schlag alle transfiniten Zahlen α zu definieren! Diesem Ziel nähern wir uns in mehreren Schritten. Und erst in Kapitel 6 gelangen wir schließlich zu einer Definition der Ordinalzahlen, werden dann aber bereits alle wesentlichen Resultate über die Struktur der Reihe (+) zur Verfügung haben.

Unsere Intuition bringt noch keine konkreten Vorstellungen über die Natur der in der Reihe vorkommenden Objekte α mit sich, von einem Anfangsstück abgesehen, das wir mit den natürlichen Zahlen, dem neuen Zeichen ω und arithmetischen Ausdrücken wie $\omega + 1$, $\omega + \omega$, usw. bilden können. Liegen nun die Stufen α der Reihe als Ganzes weitgehend im Nebel, so ist doch das Wandern auf ihnen beschreibbar: Wir besitzen eine gute Intuition über den Verlauf der Reihe, das Wesen der ihr innewohnenden Ordnung, und wir versuchen nun zunächst, die charakteristischen Eigenschaften dieser Ordnung herauszufinden.

(W1) Je zwei verschiedene Elemente α und α' der Reihe sind vergleichbar in dem Sinne, daß α vor α' oder aber α' vor α in der Reihe erscheint.
Wir schreiben $\alpha < \alpha'$, falls α vor α' in der Reihe erscheint.

(W2) Die Reihe hat ein erstes Element, das wir mit 0 bezeichnen.

(W3) Jedes α hat einen eindeutigen Nachfolger β , den wir auch durch $\alpha + 1$ bezeichnen. Wir können den Nachfolger β von α auch so charakterisieren:

$\beta =$ „das kleinste Element β' der Reihe, für das $\alpha < \beta'$ gilt“.

$\alpha + 1$ schließt unmittelbar an α an:

$$0, 1, \dots, \dots, \dots, \alpha, \alpha + 1, \dots, \dots, \dots$$

(W4) Jedes Anfangsstück A der Reihe hat einen eindeutigen Nachfolger $\beta = \beta(A)$. Dieses β ist charakterisiert durch:

$\beta =$ „das kleinste Element β' der Reihe für das gilt: $\alpha < \beta'$ für alle $\alpha \in A$ “.

$\beta(A)$ schließt unmittelbar an A an:

$$\text{--- Anfangsstück A ---, } \beta(A), \dots, \dots, \dots$$

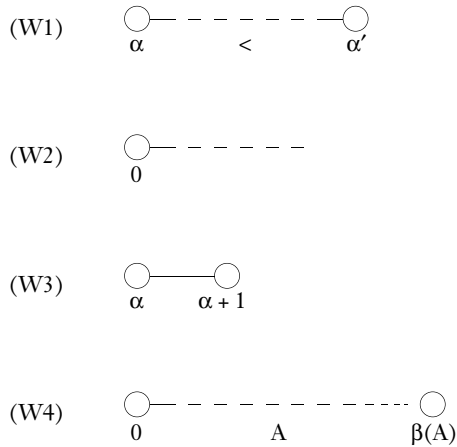
Ist z. B. $A = \{0, 1, 2, \dots\} = \mathbb{N}$, so ist $\beta(A) = \omega$.

Ist $A = \{0, 1, 2, \dots, \omega, \dots, \alpha\}$, so ist $\beta(A) = \alpha + 1$.

Diese vier Eigenschaften sind entscheidend! Mit ihrer Hilfe können wir zwar noch keine unmittelbare Definition der Ordinalzahlen selbst geben, jedoch gewinnen wir aus ihnen eine Form, in die genau die Wegstrecken hineinpassen, auf denen ein transfiniter Prozess abläuft. Cantor hat für diese Form den Begriff der Wohlordnung geprägt. Die Idee ist, daß Wohlordnungen genau wie Anfangsstücke der Reihe (+) aussehen, wenn man von den Namen der verwendeten Objekte, etwa 0, 1,

2, ..., ω , usw. absieht, und nur noch die unter ihnen herrschende Ordnung betrachtet. Verstehen wir den Verlauf von Anfangsstücken der Reihe (+), also ihre Struktur bis zu einer gedachten Stelle α , so wird es leichter möglich sein, die Stelle α selbst, die Ordinalzahl α , als Objekt zu definieren.

Die Form der Wohlordnung besteht aus vier Bedingungen, die nur dahingehend von (W1) – (W4) abweichen, daß Anfangsstücke der Reihe (+) irgendwann eine Ende haben, im Gegensatz zur gesamten Reihe selbst. Cantor definiert Wohlordnungen wie folgt.



Cantor (1883b): „Ein anderer großer, den neuen Zahlen zuzuschreibender Gewinn besteht für mich in einem *neuen*, bisher noch nicht vorgekommenen Begriffe der *Anzahl* der Elemente einer *wohlgeordneten* unendlichen Mannigfaltigkeit ...

Unter einer *wohlgeordneten* Menge ist jede wohldefinierte Menge zu verstehen, bei welcher die Elemente durch eine bestimmt vorgegebene Sukzession mit einander verbunden sind [(W1)], welcher gemäß es ein *erstes* Element der Menge gibt [(W2)] und sowohl auf jedes einzelne Element (falls es nicht das letzte in der Sukzession ist) ein bestimmtes anderes folgt [(W3')], wie auch zu jeder beliebigen endlichen oder unendlichen Menge von Elementen ein bestimmtes Element gehört, welches das ihnen allen *nächst* folgende Element in der Sukzession ist (es sei denn, daß es ein ihnen allen in der Sukzession folgendes überhaupt nicht gibt) [(W4')].“

Diese Definition gibt die zugrunde liegenden Ordnungseigenschaften am direktesten wieder. Sie ist aber etwas schwerfällig, und es gibt eine griffigere äquivalente Umformulierung.

Zunächst formulieren wir die Bedingung (W1) genauer.

Definition (*lineare Ordnung*)

Sei M eine Menge, und sei $< \subseteq M \times M$ eine zweistellige Relation auf M .
 $<$ heißt eine *lineare Ordnung auf M* , falls für alle $a, b, c \in M$ gilt:

- | | |
|---|---------------------------|
| (i) $\text{non}(a < a)$, | (Irreflexivität von $<$) |
| (ii) $a < b$ und $b < c$ folgt $a < c$, | (Transitivität von $<$) |
| (iii) $a < b$ oder $a = b$ oder $b < a$. | (Linearität von $<$) |

Verbunden mit diesem zentralen Begriff sind einige Schreib- und Sprechweisen, die den Umgang mit linearen Ordnungen erleichtern. Ihre Einführung ist katalogartig und etwas langatmig, zumal sich die meisten Dinge fast von selbst verstehen.

- (1) Wir schreiben wie üblich $a < b$ für $(a, b) \in <$.
 Weiter schreiben wir $a \leq b$ für „ $a < b$ oder $a = b$ “.
- (2) Ist $X \subseteq M$ und $s \in M$, so bedeutet $X < s$, daß $x < s$ für alle $x \in X$ gilt.
 Analog sind $X \leq s$, $s < X$, $s \leq X$ definiert.
 s heißt *<-kleinstes Element der Ordnung*, falls $s \leq M$. Analog ist *<-größtes Element* definiert.
- (3) Ist $<$ eine lineare Ordnung auf M , so nennen wir auch das Paar $(M, <)$ eine lineare Ordnung. Wir verwenden auch die Schreibweise $\langle M, < \rangle$ für $(M, <)$, und nennen $\langle M, < \rangle$ eine (Ordnungs-) *Struktur*. Die Menge M heißt der *Träger* der Ordnung $\langle M, < \rangle$.
- (4) Wir notieren lineare Ordnungen $\langle M, <_M \rangle$, $\langle N, <_N \rangle$ oft einfach als $\langle M, < \rangle$ und $\langle N, < \rangle$, wobei die beiden $<$ -Relationen i. a. nichts miteinander zu tun haben. Diese ungefährliche Konvention erleichtert die Lesbarkeit.
 (Cantor hat die Erwähnung der Ordnungen oft ganz unterdrückt.)
- (5) Ist $\langle M, < \rangle$ eine lineare Ordnung und $N \subseteq M$, so sei $<|N = < \cap (N \times N)$ die *Einschränkung der Ordnung auf N* . Dann ist $\langle N, <|N \rangle$ eine lineare Ordnung. Wir schreiben für derartige Ordnungen auch kurz $\langle N, < \rangle$, und meinen mit $<$ dann die Einschränkung von $<$ auf N .

Beispiele für lineare Ordnung sind $\langle \mathbb{N}, < \rangle$, $\langle \mathbb{Z}, < \rangle$, $\langle \mathbb{Q}, < \rangle$, $\langle \mathbb{R}, < \rangle$ mit den üblichen Ordnungen auf den Zahlen. $\langle \emptyset, \emptyset \rangle$ ist die triviale lineare Ordnung, und für jedes Objekt x ist $\langle \{x\}, \emptyset \rangle$ eine lineare Ordnung mit einem einelementigen Träger.

Cantor (1895): „Eine Menge M nennen wir ‚*einfach geordnet*‘ [linear geordnet], wenn unter ihren Elementen m eine bestimmte ‚*Rangordnung*‘ herrscht, in welcher von je zwei beliebigen Elementen m_1 und m_2 das eine den ‚*niedrigeren*‘, das andere den ‚*höheren*‘ Rang einnimmt, und zwar so, daß wenn von drei Elementen m_1 , m_2 und m_3 etwa m_1 dem Range nach niedriger ist als m_2 , dieses niedriger als m_3 , alsdann auch immer m_1 niedrigeren Rang hat als m_3 .

Die Beziehung zweier Elemente m_1 und m_2 , bei welcher m_1 den niedrigeren, m_2 den höheren Rang in der gegebenen Rangordnung hat, soll durch die Formeln ausgedrückt werden

$$(1) \quad m_1 < m_2, \quad m_2 > m_1. "$$

Übung

Sei M eine Menge, und sei $< = \subset \mid M = \subset \cap M \times M$, also

$a < b$ gdw $a, b \in \mathcal{P}(M)$ und $a \subset b$.

Dann ist $<$ irreflexiv und transitiv auf M , aber i. a. nicht linear.

Man nennt irreflexive und transitive Relationen $< \subseteq M \times M$ *partielle Ordnungen auf M* . Lineare und allgemeiner partielle Ordnungen spielen in der Mengenlehre eine große Rolle und haben eine reiche Theorie. In den „Grundzügen der Mengenlehre“ von Hausdorff (1914) werden Ordnungen bereits sehr ausführlich und allgemein untersucht (Kapitel 4 – 6). Hier interessieren uns zunächst nur die Ordnungen, die zu den Ordinalzahlen führen, und später dann die Ordnungen der rationalen und der reellen Zahlen.

Damit ist die Bedingung (W1) scharf gefaßt. Die drei übrigen Bedingungen können wir nun zu einer einzigen Bedingung verschmelzen: Der Nachfolger eines Elementes oder einer Teilmenge ist jeweils das *Minimum* aller größeren Elemente, und das erste Element ist das *Minimum* der ganzen Ordnung. Entscheidend ist die Existenz dieses Minimums für jede nichtleere Teilmenge – eine vertraute Eigenschaft, wenn man an die natürlichen Zahlen denkt.

Definition (*Wohlordnung; zweite Fundamentaldefinition der Mengenlehre*)

Sei $\langle M, < \rangle$ eine lineare Ordnung.

$\langle M, < \rangle$ heißt eine *Wohlordnung auf M* , falls jede nichtleere Teilmenge M' von M ein $<$ -kleinstes Element besitzt, d. h. es gibt ein $m \in M$ mit:

- (i) $m \in M'$,
- (ii) $m \leq x$ für alle $x \in M'$.

Die Forderung $m \in M'$ ist wesentlich. Die reellen Zahlen $\langle \mathbb{R}, < \rangle$ sind z. B. keine Wohlordnung: Für $M' = \{1/n \mid n \in \mathbb{N}, n > 0\} \subseteq \mathbb{R}$ existiert $0 = \inf(M')$ in $\mathbb{R}$, aber es ist $0 \notin M'$; M' hat kein $<$ -kleinstes Element.

Nichtleere Teilmengen von Wohlordnungen haben i. a. kein größtes Element: $\mathbb{N} \subseteq M$ hat kein größtes Element in $M = \{0, 1, 2, \dots, \omega\}$, versehen mit der natürlichen Ordnung.

Die triviale Struktur $\langle \emptyset, \emptyset \rangle$ ist nach dieser Definition eine Wohlordnung, eine Konvention, die vielfach nützlich ist, und dem Status der 0 als natürlicher Zahl entspricht. Für jedes $n \in \mathbb{N}$ ist $\langle \bar{n}, < \rangle = \langle \{0, 1, \dots, n-1\}, < \rangle$ eine Wohlordnung mit der üblichen $<$ -Relation. Weiter ist z. B. die im ersten Kapitel betrachtete Teilmenge $M = \{0\} \cup \{n-1/k \mid n, k \in \mathbb{N}, n \geq 1, k \geq 2\}$ der reellen Zahlen wohlgeordnet unter der üblichen Ordnung auf $\mathbb{R}$.

Die so definierten Wohlordnungen sind nun in der Tat genau die linearen Ordnungen, welche die von Cantor genannten Bedingungen erfüllen – wobei wir hier die leere Menge als wohlgeordnet ansehen, und daher auch die Bedingung der Existenz eines kleinsten Elementes modifizieren müssen:

Übung

Zeigen Sie, daß die Wohlordnungen genau die Wohlordnungen im Sinne der Definition von Cantor sind. Zu zeigen ist also, daß die Aussagen (A) und (B) äquivalent sind, wobei:

- (A) $<$ ist eine Wohlordnung auf M (nach obiger Definition).
- (B) (W1) $<$ ist eine lineare Ordnung auf M .
 (W2') Ist $M \neq \emptyset$, so existiert ein $<$ -kleinstes Element von M .
 (W3') Ist $\alpha \in M$ und existiert überhaupt ein $\beta \in M$ mit $\alpha < \beta$, so existiert ein $<$ -kleinstes β mit $\alpha < \beta$.
 (W4') Ist $A \subseteq M$ und existiert überhaupt ein $\beta \in M$ mit $A < \beta$, so existiert ein $<$ -kleinstes $\beta = \beta(A) \in M$ mit $A < \beta$.

Cantor (1897):

„A. „Jede [nichtleere] Teilmenge F_1 einer [nach Cantors Definition (W1) – (W4')] wohlgeordneten Menge F hat ein niederstes Element.“ ...

B. „Ist eine einfach [linear] geordnete Menge F so beschaffen, daß sowohl F , wie auch jede ihrer [nichtleeren] Teilmengen ein niederstes Element hat, so ist F eine wohlgeordnete Menge.““

Damit ist das Fundament der transfiniten Prozesse ausgegraben. Sie verlaufen entlang Wohlordnungen, also linearen Ordnungen, in denen nichtleere Teilmengen immer ein kleinstes Element aufweisen. Der Verlauf eines solchen Prozesses entlang einer Wohlordnung ist in der Cantorsche Definition besonders anschaulich. Die Elemente der Wohlordnung markieren die Stufen des Prozesses.

In Wohlordnungen zerfallen die Elemente in zwei Typen, je nachdem, ob ein unmittelbares Vorgängerelement existiert oder nicht. Zudem hat das erste Element der Wohlordnung eine besondere Stellung.

Definition (*Nachfolgerelement und Limeselement einer Wohlordnung*)

Sei $\langle M, < \rangle$ eine Wohlordnung und sei $x \in M$. x heißt:

- (i) *Anfangselement* von $\langle M, < \rangle$, falls x das $<$ -kleinste Element von M ist.
 (ii) *Nachfolgerelement* von $\langle M, < \rangle$, falls ein $y \in M$ existiert mit $x = \text{„das } <\text{-kleinste } x' \in M \text{ mit } y < x'\text{“}$.

Wir schreiben $x = y + 1$ und $y = x - 1$ in diesem Fall.

x heißt dann der *Nachfolger* von y , y der *Vorgänger* von x in $\langle M, < \rangle$.

- (iii) *Limeselement* von $\langle M, < \rangle$, falls x kein Nachfolgerelement und nicht das Anfangselement von $\langle M, < \rangle$ ist.

Einfache Operationen mit Wohlordnungen

Wir besprechen noch einige Möglichkeiten, aus gegebenen Wohlordnungen neue zu konstruieren.

Eine fast triviale, aber wichtige Beobachtung ist, daß jede Teilmenge einer Wohlordnung durch die Ordnung der Ausgangsmenge wiederum wohlgeordnet wird:

Übung

Sei $\langle M, < \rangle$ eine Wohlordnung, und sei $N \subseteq M$.
Dann ist $\langle N, < \rangle = \langle N, <|_N \rangle$ eine Wohlordnung.

Ausgezeichnete Teilmengen sind die Anfangsstücke einer Wohlordnung:

Definition (*Anfangsstück einer Wohlordnung*)

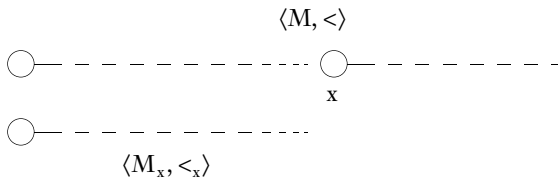
Sei $\langle M, < \rangle$ eine Wohlordnung, und sei $x \in M$. Wir setzen:

$$M_x = \{ y \in M \mid y < x \},$$

$$<_x = <|_{M_x}, \text{ d.h. } <_x = < \cap (M_x \times M_x).$$

$\langle M_x, <_x \rangle$ heißt das *durch x bestimmte Anfangsstück* von $\langle M, < \rangle$.

Eine Wohlordnung $\langle N, < \rangle$ heißt ein (*echtes*) *Anfangsstück* von $\langle M, < \rangle$, falls ein $x \in M$ existiert mit $\langle N, < \rangle = \langle M_x, <_x \rangle$.



Jedes $\langle M_x, <_x \rangle$ ist eine Wohlordnung und stimmt bis x mit der Wohlordnung $\langle M, < \rangle$ überein. Ist x das Anfangselement von $\langle M, < \rangle$, so ist $\langle M_x, <_x \rangle = \langle \emptyset, \emptyset \rangle$. Für alle $x \in M$ und alle $y \in M_x$ ist $(M_x)_y = M_y$.

Umgekehrt kann man Anfangsstücke zu einer Wohlordnung vereinigen:

Definition (*Vereinigung von vergleichbaren Wohlordnungen*)

Sei Γ eine Menge von Wohlordnungen mit der Eigenschaft:

Sind $\langle A, < \rangle$ und $\langle B, < \rangle$ verschiedene Elemente von Γ , so ist
 $\langle A, < \rangle$ ein Anfangsstück von $\langle B, < \rangle$ oder $\langle B, < \rangle$ ein Anfangsstück von $\langle A, < \rangle$.

Wir setzen $\bigcup \Gamma = \langle N_\Gamma, <_\Gamma \rangle$, wobei

$$N_\Gamma = \bigcup \{ M \mid \text{es existiert ein } < \text{ mit } \langle M, < \rangle \in \Gamma \},$$

$$<_\Gamma = \bigcup \{ < \mid \text{es existiert ein } M \text{ mit } \langle M, < \rangle \in \Gamma \}.$$

Übung

Sei Γ wie in der Definition.

Dann ist $\langle \mathbb{N}, < \rangle = \bigcup \Gamma$ eine Wohlordnung.

Ist $x \in \mathbb{N}$ und $x \in M$ für ein $\langle M, < \rangle \in \Gamma$, so ist $\langle \mathbb{N}_x, <_x \rangle = \langle M_x, <_x \rangle$.

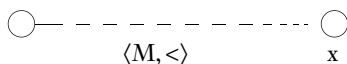
Weiter lassen sich Wohlordnungen durch Aneinanderhängen zu neuen Wohlordnungen kombinieren. Wir besprechen die Arithmetik linearer Ordnungen allgemein in Kapitel 8, sodaß wir uns hier auf einen einfachen Spezialfall beschränken können, der im folgenden eine Rolle spielt.

Definition (*Enderweiterung einer Wohlordnung um ein Element*)

Sei $\langle M, < \rangle$ eine Wohlordnung, und sei $x \notin M$.

Dann ist die *Enderweiterung von $\langle M, < \rangle$ um x* , in Zeichen $\langle M, < \rangle + \{x\}$, definiert durch:

$$\langle M, < \rangle + \{x\} = \langle M \cup \{x\}, < \cup (M \times \{x\}) \rangle.$$



$\langle M, < \rangle + \{x\}$ ist eine Wohlordnung, und es gilt $y < x$ für alle $y \in M$, d. h. x ist das größte Element der erweiterten Wohlordnung. So ist z. B. $\langle \{0, 1, \dots, \omega\}, < \rangle = \langle \mathbb{N}, < \rangle + \{\omega\}$.

In den beiden nächsten Kapiteln untersuchen wir den Begriff der Wohlordnung im Detail. Zunächst klären wir das Verhältnis verschiedener Wohlordnungen untereinander, und zeigen, daß die *Länge* einer Wohlordnung bereits vollständig charakterisiert. Danach behandeln wir die Frage nach der Existenz von Wohlordnungen auf beliebigen Mengen. Dieses Problem war im 19. Jahrhundert offengeblieben, wurde dann aber durch Ernst Zermelo 1904 positiv beantwortet. Im sechsten Kapitel gewinnen wir schließlich aus dem Wohlordnungsbegriff den Begriff der Ordinalzahl: Zunächst klassisch in „edler Einfalt und stiller Größe“ als *Ordnungstypus*, als das allen Wohlordnungen gleicher Länge Gemeinsame, und dann modern über bestimmte, besonders ausgezeichnete Wohlordnungen.

Felix Hausdorffs Einführung des Wohlordnungsbegriffs

„Bei dem Versuch (Kap. IV, § 1), die Eigenschaften der natürlichen Zahlenreihe auf unendliche Mengen zu übertragen, haben wir zunächst das Moment der Ordnung berücksichtigt. Die Zahlenreihe ist aber eine sehr spezielle geordnete Menge, und ihre Funktion als Instrument zum Zählen knüpft sich gerade an eine solche spezielle Eigenschaft, daß nämlich, wenn man bis n gezählt hat, nunmehr eine nächstfolgende Zahl $n + 1$ an die Reihe kommt. Anders ausgedrückt: wir haben hier eine geordnete Menge A von der Beschaffenheit, daß bei jeder Zerlegung $A = P + Q$ das Endstück Q (falls es Elemente enthält) ein erstes Element hat. Diese Eigenschaft übertragen wir auf beliebige Mengen. Eine [linear]

geordnete Menge A heie wohlgeordnet und ihr Ordnungstypus eine Ordnungszahl, wenn jedes von Null verschiedene Endstck ein erstes Element hat.

Wir knnen auch sagen: A ist wohlgeordnet, wenn jede von Null verschiedene Teilmenge A' ein erstes Element hat. Diese Bedingung ist ja offenbar hinreichend, aber auch notwendig; denn (Kap. IV, § 4) A' bestimmt das mit ihm koinitiale Endstck $[= \{x \in A \mid \text{es gibt ein } y \in A' \text{ mit } y \leq x\}]$, dessen erstes Element auch das erste Element von A' ist. Es ist in der Definition eingeschlossen, da auch A selbst ein erstes Element haben mu ...“

(Felix Hausdorff 1914, „Grundzge der Mengenlehre“)

4. Der Fundamentalsatz über Wohlordnungen

Gehen wir in einer Wohlordnung von links nach rechts, vom Anfangselement zu immer größeren Elementen, so bekommen wir schnell den Eindruck, daß der Weg, so weit er führt, determiniert ist: Wir können entweder einen Nachfolgerschritt tun, oder zu einem Limeselement springen. Wohlordnungen scheinen also durch ihre „Länge“ eindeutig festgelegt zu sein, wenn man von der Natur der Elemente des Trägers der Wohlordnung absieht.

Dies wollen wir nun beweisen. Der Leser vergleiche unser Vorgehen und die Resultate mit der Diskussion des Vergleichs zweier Mengen nach ihrer Mächtigkeit. Hier vergleichen wir nun zwei Wohlordnungen nach ihrer Länge.

Längenvergleiche

Zunächst definieren wir, wann zwei Wohlordnungen gleichlang sind. Cantor hält nach seiner ersten Definition des Wohlordnungsbegriffs sogleich fest:

Cantor (1883b): „Zwei ‚wohlgeordnete‘ Mengen werden nun von derselben *Anzahl* (mit Bezug auf die für sie vorgegebenen Sukzessionen) genannt, wenn eine gegenseitig eindeutige Zuordnung [Bijektion] derselben derart möglich ist, daß, wenn E und F irgend zwei Elemente der einen, E_1 und F_1 die entsprechenden Elemente der anderen sind, immer die Stellung von E und F in der Sukzession der ersten Menge in Übereinstimmung ist mit der Stellung von E_1 und F_1 in der Sukzession der zweiten Menge ...“

In heutiger Sprache:

Definition (*Wohlordnungen gleicher Länge oder ähnliche Wohlordnungen*)

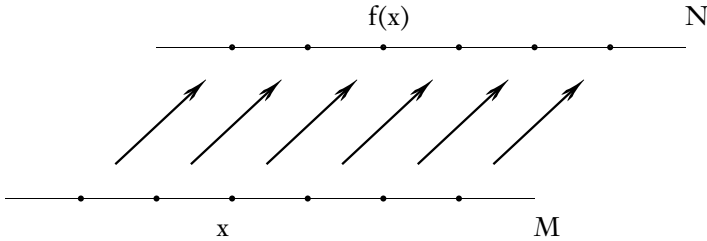
Zwei Wohlordnungen $\langle M, < \rangle$ und $\langle N, < \rangle$ heißen *gleichlang* oder *ähnlich*, falls eine Bijektion $f : M \rightarrow N$ existiert derart, daß für alle $x, y \in M$ gilt:

$$x < y \text{ gdw } f(x) < f(y).$$

Ein solches f heißt ein *Ordnungsisomorphismus* zwischen $\langle M, < \rangle$ und $\langle N, < \rangle$.

Wir schreiben $\langle M, < \rangle \equiv \langle N, < \rangle$, falls $\langle M, < \rangle$ und $\langle N, < \rangle$ gleichlang sind.

In seinen späteren Arbeiten benutzt Cantor „ähnlich“ (1895, 1897). Heute ist darüber hinaus auch „ordnungsisomorph“ gebräuchlich.



Übung

- (a) $\equiv$ ist eine Äquivalenzrelation auf der Klasse der Wohlordnungen.
 (b) Seien $\langle M, < \rangle$ und $\langle N, < \rangle$ Wohlordnungen. Weiter sei $f : M \rightarrow N$ surjektiv und für alle $x, y \in M$ gelte: $x < y$ folgt $f(x) < f(y)$.
 Zeigen Sie, daß f ein Ordnungsisomorphismus ist.

Man kann andererseits strukturerhaltende Abbildungen benutzen, um von einer Struktur zu zeigen, daß sie eine Wohlordnung ist:

Übung

Seien $\langle M, < \rangle$ eine Wohlordnung, N eine Menge, und sei R eine zweistellige Relation auf N . Weiter sei $f : M \rightarrow N$ bijektiv mit:

Für alle $x, y \in M$ gilt: $x < y$ gdw $f(x) R f(y)$.

Dann ist $\langle N, R \rangle$ eine Wohlordnung mit $\langle M, < \rangle \equiv \langle N, R \rangle$.

Sind $\langle M, < \rangle$ und $\langle N, < \rangle$ gleichlang, so ist wegen der Bijektivität eines Ordnungsisomorphismus insbesondere $|M| = |N|$. Umgekehrt kann man Wohlordnungen durch Bijektionen übertragen:

Übung (induzierte Wohlordnung)

Seien $\langle M, < \rangle$ eine Wohlordnung und N eine Menge mit $|M| = |N|$.

Weiter sei $f : M \rightarrow N$ bijektiv. Für $x, y \in N$ setzen wir:

$x < y$ falls $f^{-1}(x) < f^{-1}(y)$.

Dann ist $\langle N, < \rangle$ eine Wohlordnung mit $\langle M, < \rangle \equiv \langle N, < \rangle$.

Ähnliche Wohlordnungen $\langle M, < \rangle$ und $\langle N, < \rangle$ kann man wie folgt deuten: Wir haben die gleiche Ordnung vor uns, lediglich werden die Elemente der Ordnung anders benannt. Ein Ordnungsisomorphismus $f : M \rightarrow N$ übersetzt einfach die Namen der Elemente. Wir werden gleich sehen, daß eine solche Funktion f eindeutig bestimmt ist.

Ein Ordnungsisomorphismus f erhält alle Eigenschaften, die sich mit „<“ formulieren lassen; ist z. B. y der Nachfolger von x in M , so ist $f(y)$ der Nachfolger von $f(x)$ in N ; ist z ein Limeselement in N , so ist $f^{-1}(z)$ ein Limeselement in M , u. s. w. (Eine präzise Formulierung und einen allgemeinen Beweis dieser Übertragung von Eigenschaften durch einen Isomorphismus kann man in der *Modelltheorie* angeben.)

Mit Hilfe des Begriffs eines Anfangsstücks können wir nun aus „gleichlang“ leicht auch „kürzer als“ für Wohlordnungen definieren.

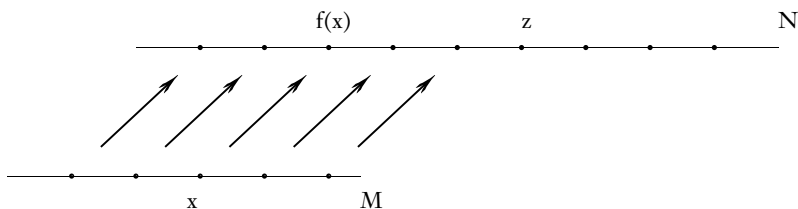
Definition (die Relation „kürzer als“ für Wohlordnungen)

Seien $\langle M, < \rangle$ und $\langle N, < \rangle$ Wohlordnungen.

$\langle M, < \rangle$ heißt *kürzer als* $\langle N, < \rangle$, falls ein $z \in N$ existiert mit der Eigenschaft:

$\langle M, < \rangle$ und $\langle N_z, <_z \rangle$ sind gleichlang.

Wir schreiben $\langle M, < \rangle \triangleleft \langle N, < \rangle$, falls $\langle M, < \rangle$ kürzer als $\langle N, < \rangle$ ist.



Eine Wohlordnung ist also kürzer als eine andere, wenn sie nach einer Umbenennung ihrer Elemente ein Anfangsstück der anderen darstellt.

Übung

- (i) $\triangleleft$ ist eine transitive Relation auf der Klasse der Wohlordnungen.
- (ii) Ist $\langle M, < \rangle \triangleleft \langle N, < \rangle$ und $\langle N, < \rangle \equiv \langle Q, < \rangle$, so ist $\langle M, < \rangle \triangleleft \langle Q, < \rangle$.
- (iii) Ist $\langle M, < \rangle \triangleleft \langle N, < \rangle$ und $\langle M, < \rangle \equiv \langle Q, < \rangle$, so ist $\langle Q, < \rangle \triangleleft \langle N, < \rangle$.

Für die Irreflexivität, d.h. für $\text{non}(\langle M, < \rangle \triangleleft \langle M, < \rangle)$, ist ein kleiner Trick nötig.

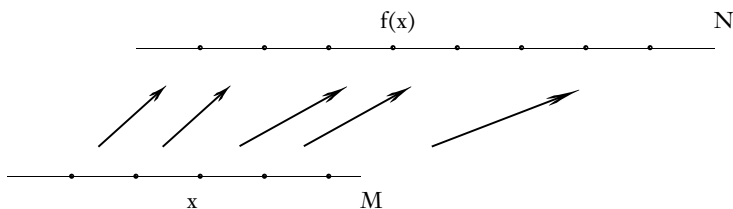
Ordnungstreue Abbildungen

Eine ordnungstreue Abbildung zwischen zwei Wohlordnungen ist eine Funktion, die die Lage von je zwei Elementen untereinander erhält.

Definition (ordnungstreue Abbildungen)

Seien $\langle M, < \rangle$ und $\langle N, < \rangle$ Wohlordnungen, und sei $f : M \rightarrow N$ eine Funktion. f heißt *ordnungstreu*, falls für alle $x, y \in M$ gilt:

$x < y$ gdw $f(x) < f(y)$.



Insbesondere ist ein ordnungstreuere f injektiv. Ist f bijektiv, so ist f ein Ordnungsisomorphismus, und $\langle M, < \rangle$ und $\langle N, < \rangle$ sind gleichlang. Für „ordnungstreu“ genügt bereits die Bedingung „ $x < y$ folgt $f(x) < f(y)$ “, wie man leicht sieht.

Der Trick zum Nachweis der Irreflexivität von $\triangleleft$ besteht nun in folgender Beobachtung, die auf Ernst Zermelo zurückgeht: Eine ordnungstreuere Abbildung einer Wohlordnung in sich selbst kann Elemente nur nach oben verschieben.

Satz (ordnungstreuere $f: M \rightarrow M$ sind expansiv)

Sei $\langle M, < \rangle$ eine Wohlordnung, und sei $f: M \rightarrow M$ ordnungstreu.

Dann gilt $x \leq f(x)$ für alle $x \in M$.

Beweis

Annahme nicht. Dann ist $Y = \{x \in M \mid f(x) < x\} \neq \emptyset$.

Sei y das kleinste Element von Y . Dann gilt $y \in Y$, also $f(y) < y$.

Wegen f ordnungstreu ist dann aber $f(f(y)) < f(y)$. Also ist auch $f(y) \in Y$.

– Aber $f(y) < y$, im Widerspruch zur Minimalität von y .

Zermelo (1932, Anmerkung zu § 13 von Cantor 1897): „Die Sätze A – M dieses Paragraphen sind größtenteils lediglich Hilfssätze zum Beweis des ‚Ähnlichkeitssatzes‘ [Vergleichbarkeitssatzes für Wohlordnungen] ..., in welchem die elementare Theorie der wohlgeordneten Mengen gipfelt. Hier lassen sich aber die Sätze B – F einfacher als bei Cantor beweisen bzw. ersetzen durch Voranstellung des allgemeinen (vom Herausgeber [Zermelo] herrührenden) Hilfssatzes: ‚Bei keiner ähnlichen Abbildung einer wohlgeordneten Menge auf einen ihrer Teile wird ein Element a auf ein vorangehendes $a' < a$ abgebildet‘ ...“

Korollar (Eindeutigkeit des Ordnungsisomorphismus)

(i) Sei $\langle M, < \rangle$ eine Wohlordnung.

Dann ist $f = \text{id}_M$ der einzige Ordnungsisomorphismus $f: M \rightarrow M$.

(ii) Seien $\langle M, < \rangle$ und $\langle N, < \rangle$ gleichlang.

Dann existiert genau ein Ordnungsisomorphismus $f: M \rightarrow N$.

Beweis

zu (i): Offenbar ist $\text{id}_M: M \rightarrow M$ ordnungsisomorph.

Ist $f: M \rightarrow M$ ordnungsisomorph, so sind $f: M \rightarrow M$ und $f^{-1}: M \rightarrow M$ ordnungstreu. Also gilt $f(x) \geq x$ und $f^{-1}(x) \geq x$ für alle $x \in M$.

Annahme, $f(x) > x$ für ein $x \in M$.

Dann gilt aber $x = f^{-1}(f(x)) \geq f(x) > x$, Widerspruch!

Also gilt $f(x) = x$ für alle $x \in M$.

zu (ii): Seien $f: M \rightarrow N$ und $g: M \rightarrow N$ ordnungsisomorph.

Dann ist $g^{-1} \circ f: M \rightarrow M$ ordnungsisomorph, also $g^{-1} \circ f = \text{id}_M$,

– und damit $f = g$.

Eine weitere Folgerung ist, daß Teilmengen einer Wohlordnung nicht länger sein können als die Wohlordnung selbst, und daß verschiedene Anfangsstücke einer Wohlordnung verschiedene Längen haben:

Korollar (*Irreflexivität von $\triangleleft$ und Nichtäquivalenz der Anfangsstücke*)

Sei $\langle M, < \rangle$ eine Wohlordnung. Dann gilt:

- (i) Für alle $N \subseteq M$ gilt $\text{non}(\langle M, < \rangle \triangleleft \langle N, < \rangle)$.
Insbesondere also $\text{non}(\langle M, < \rangle \triangleleft \langle M, < \rangle)$.
- (ii) Für alle $x \in M$ gilt $\langle M_x, < \rangle \neq \langle M, < \rangle$.
- (iii) Für alle $x, y \in M$ mit $x \neq y$ ist $\langle M_x, < \rangle \neq \langle M_y, < \rangle$.

Beweis

zu (i): *Annahme* $\langle M, < \rangle \triangleleft \langle N, < \rangle$ für ein $N \subseteq M$. Dann existieren ein $x \in N$ und ein Ordnungsisomorphismus $f: M \rightarrow N_x$. Wegen $f(x) \in N_x$ ist dann $f(x) < x$. Aber $f: M \rightarrow M$ ist ordnungstreu, also $x \leq f(x)$, *Widerspruch!*

zu (ii): *Annahme* es gibt ein $x \in M$ mit $\langle M_x, < \rangle \equiv \langle M, < \rangle$.

Dann gilt $\langle M_x, < \rangle \triangleleft \langle M, < \rangle \equiv \langle M_x, < \rangle$, *im Widerspruch* zu (i).

zu (iii): O.E. ist $x < y$. Dann ist $\langle M_x, < \rangle$ ein Anfangsstück von $\langle M_y, < \rangle$.

— Dann aber $\langle M_x, < \rangle \neq \langle M_y, < \rangle$ nach (ii).

Aus dem Fundamentalsatz wird folgen, daß eine Teilordnung einer Wohlordnung kürzer oder gleichlang zur ursprünglichen Wohlordnung ist.

Der Fundamentalsatz

Wir zeigen nun den „Vergleichbarkeitssatz für Wohlordnungen“ von Georg Cantor, vollständig bewiesen im Jahre 1897: Zwei Wohlordnungen sind gleichlang, oder die eine ist kürzer als die andere. Anders formuliert: $\triangleleft$ ist eine lineare Ordnung auf der Klasse der Wohlordnungen modulo der Äquivalenz „gleichlang“.

Fraenkel (1928): „Bevor wir weitere Eigenschaften der wohlgeordneten Mengen entwickeln ..., wollen wir sogleich den Vorzug der wohlgeordneten Mengen kennenlernen, der innerhalb der allgemeinen Theorie im Mittelpunkt steht und die überragende Bedeutung der wohlgeordneten Mengen begründet: den Vorzug der Vergleichbarkeit.“

Satz (*Vergleichbarkeitssatz für Wohlordnungen*)

Seien $\langle M, < \rangle$ und $\langle N, < \rangle$ zwei Wohlordnungen.

Dann gilt genau einer der drei Fälle:

$\langle M, < \rangle \triangleleft \langle N, < \rangle$ oder $\langle M, < \rangle \equiv \langle N, < \rangle$ oder $\langle N, < \rangle \triangleleft \langle M, < \rangle$.

Beweis

Wegen der Irreflexivität und Transitivität von $\triangleleft$ tritt allenfalls einer der drei Fälle ein. Wir setzen:

$$f = \{ (x, y) \mid x \in M, y \in N, \langle M_x, < \rangle \equiv \langle N_y, < \rangle \}.$$

Dann gilt:

- (i) f ist eine injektive Funktion,
- (ii) $x \in \text{dom}(f), z < x$ folgt $z \in \text{dom}(f)$ für alle $x, z \in M$,
- (iii) $y \in \text{rng}(f), z < y$ folgt $z \in \text{rng}(f)$ für alle $y, z \in N$,
- (iv) $f : \text{dom}(f) \rightarrow \text{rng}(f)$ ist ordnungstreu,
- (v) $\text{dom}(f) = M$ oder $\text{rng}(f) = N$.

Beweis hierzu

zu (i): Seien $(x, y_1), (x, y_2) \in f$. Dann ist $\langle M_x, < \rangle \equiv \langle N_{y_1}, < \rangle$ und $\langle M_x, < \rangle \equiv \langle N_{y_2}, < \rangle$. Damit gilt $\langle N_{y_1}, < \rangle \equiv \langle N_{y_2}, < \rangle$, also $y_1 = y_2$. Also ist f eine Funktion.

Analog folgt aus $(x_1, y), (x_2, y) \in f$, daß $x_1 = x_2$. Also ist f injektiv.

zu (ii): Sei $(x, y) \in f$ und sei $g : M_x \rightarrow N_y$ ordnungsisomorph. Ist $z \in M, z < x$, so ist $g|_{M_z} : M_z \rightarrow N_{g(z)}$ ordnungsisomorph. Also ist $(z, g(z)) \in f$, und damit $z \in \text{dom}(f)$.

zu (iii): Analog zu (ii): Sei $(x, y) \in f$ und sei $g : M_x \rightarrow N_y$ ordnungsisomorph. Ist $z \in N, z < y$, so ist die Funktion $g|_{M_{g^{-1}(z)}} : M_{g^{-1}(z)} \rightarrow N_z$ ordnungsisomorph. Also ist $(g^{-1}(z), z) \in f$, und damit $z \in \text{rng}(f)$.

zu (iv): Seien $(x_1, y_1), (x_2, y_2) \in f$ und sei $x_1 < x_2$. Dann gilt:
 $\langle N_{y_1}, < \rangle \equiv \langle M_{x_1}, < \rangle \triangleleft \langle M_{x_2}, < \rangle \equiv \langle N_{y_2}, < \rangle$,
 also $f(x_1) = y_1 < y_2 = f(x_2)$. Dies genügt für f ordnungstreu.

zu (v): *Annahme nicht.* Sei dann

$x =$ „das $<$ -kleinste Element von $M - \text{dom}(f)$ “,

$y =$ „das $<$ -kleinste Element von $N - \text{rng}(f)$ “.

Nach (ii, iii) ist dann $\text{dom}(f) = M_x$ und $\text{rng}(f) = N_y$.

Nach (i, iv) ist also $\langle M_x, < \rangle \equiv \langle N_y, < \rangle$.

Also $(x, y) \in f$ nach Definition von f , *Widerspruch!*

Ist $\text{dom}(f) = M$, aber $\text{rng}(f) \neq N$, so gilt $\langle M, < \rangle \triangleleft \langle N, < \rangle$.

Ist $\text{dom}(f) = M$ und $\text{rng}(f) = N$, so gilt $\langle M, < \rangle \equiv \langle N, < \rangle$.

Ist $\text{dom}(f) \neq M$, aber $\text{rng}(f) = N$, so gilt $\langle N, < \rangle \triangleleft \langle M, < \rangle$
 (betrachte $f^{-1} : N \rightarrow M$).

— Also gilt $\langle M, < \rangle \triangleleft \langle N, < \rangle$ oder $\langle M, < \rangle \equiv \langle N, < \rangle$ oder $\langle N, < \rangle \triangleleft \langle M, < \rangle$.

Die Vergleichbarkeit ist das fundamentale Resultat über Wohlordnungen. Die über die bloße Linearität der Ordnung hinausgehende Bedingung, daß nichtleere Teilmengen ein kleinstes Element haben, hat zur Folge, daß eine Wohlordnung bis auf Isomorphie durch einen einzigen natürlichen Parameter bestimmt ist, nämlich ihre Länge, und daß je zwei Wohlordnungen in ihrer Länge vergleichbar sind. Im unüberschaubaren Feld der linearen Ordnungen haben wir somit einen ausgezeichneten Pfad gefunden, den wir ohne uns zu verirren verfolgen können.

Wir halten noch einige Folgerungen aus dem Fundamentalsatz fest:

Korollar

Sei $\langle M, < \rangle$ eine Wohlordnung, und sei $N \subseteq M$.
Dann gilt $\langle N, < \rangle \triangleleft \langle M, < \rangle$ oder $\langle N, < \rangle \equiv \langle M, < \rangle$.

Beweis

- Der Fall $\langle M, < \rangle \triangleleft \langle N, < \rangle$ ist nach dem Korollar oben ausgeschlossen.

Ein $N \subset M$ ist nicht unbedingt kürzer als M selbst:

Übung

Geben Sie eine Wohlordnung $\langle M, < \rangle$ an, für die ein $N \subset M$ existiert mit $\langle N, < \rangle \equiv \langle M - N, < \rangle \equiv \langle M, < \rangle$.

Weiter erhalten wir die folgende Charakterisierung der Relation „kürzer als oder gleichlang“:

Korollar

Seien $\langle M, < \rangle, \langle N, < \rangle$ Wohlordnungen. Dann sind äquivalent:

- (i) $\langle M, < \rangle \triangleleft \langle N, < \rangle$ oder $\langle M, < \rangle \equiv \langle N, < \rangle$.
- (ii) Es existiert ein ordnungstreu $f : M \rightarrow N$.

Beweis

- (i) $\hookrightarrow$ (ii): Ein Ordnungsisomorphismus f von M auf N oder ein Anfangsstück von N ist eine ordnungserhaltende Abbildung von M nach N .
- (ii) $\hookrightarrow$ (i): Sei $f : M \rightarrow N$ ordnungstreu und $N' = \text{rng}(f) \subseteq N$.
Dann ist $\langle M, < \rangle \equiv \langle N', < \rangle$, und nach dem Korollar oben ist
– $\langle N', < \rangle \triangleleft \langle N, < \rangle$ oder $\langle N', < \rangle \equiv \langle N, < \rangle$. Dies zeigt die Behauptung.

Aus dem Korollar folgt eine Cantor-Bernstein-Version für Wohlordnungen. Setzt man: $\langle M, < \rangle \leq \langle N, < \rangle$, falls ein ordnungstreu $f : M \rightarrow N$ existiert, so zeigt das Korollar: Gilt $\langle M, < \rangle \leq \langle N, < \rangle$ und $\langle N, < \rangle \leq \langle M, < \rangle$, so gilt $\langle M, < \rangle \equiv \langle N, < \rangle$. Wir haben dieses Resultat aus dem Vergleichbarkeitssatz für Wohlordnungen gewonnen, im Gegensatz zum Satz von Cantor-Bernstein, zu dessen Beweis wir den Vergleichbarkeitssatz für Mächtigkeiten nicht benutzt haben.

Keine Beweise dieses und des vorhergehenden Kapitels verwendeten Konstruktionen der Form „ein ...“. Dies ist typisch für Wohlordnungen. Wir können immer schreiben „das $<$ -kleinste ...“, haben also direkten, definierbaren Zugriff auf Elemente von nicht-leeren Teilmengen von M . Die einzige Stelle der Wohlordnungstheorie, wo ein „ein ...“ gebraucht wird, ist der Nachweis der Existenz einer Wohlordnung auf einer beliebig gegebenen Menge. Diesen Nachweis erbringen wir im nächsten Kapitel.

Georg Cantors Beweis des Vergleichbarkeitssatzes für Wohlordnungen

„A. ,Sind zwei ähnliche wohlgeordnete Mengen F und G aufeinander [ordnungsisomorph] abgebildet, so entspricht jedem Abschnitt [Anfangsstück] A von F ein ähnlicher Abschnitt B von G , und jedem Abschnitt B von G ein ähnlicher Abschnitt A von F , und die Elemente f und g von F und G , durch welche die einander zugeordneten Abschnitte A und B bestimmt sind, entsprechen einander ebenfalls bei der Abbildung.‘ ...

B. ,Eine wohlgeordnete Menge F ist keinem ihrer Abschnitte A ähnlich.‘ ...

C. ,Eine wohlgeordnete Menge F ist keiner Teilmenge irgend eines ihrer Abschnitte A ähnlich.‘ ...

D. ,Zwei verschiedene Abschnitte A und A' einer wohlgeordneten Menge F sind nicht einander ähnlich.‘ ...

E. ,Zwei ähnliche wohlgeordnete Mengen F und G lassen sich nur auf eine einzige Weise aufeinander [ordnungsisomorph] abbilden.‘ ...

F. ,Sind F und G zwei wohlgeordnete Mengen, so kann ein Abschnitt A von F höchstens einen ihm ähnlichen Abschnitt B in G haben.‘ ...

G. ,Sind A und B ähnliche Abschnitte zweier wohlgeordneter Mengen F und G , so gibt es auch zu jedem kleineren Abschnitt $A' < A$ von F einen ähnlichen Abschnitt $B' < B$ von G und zu jedem kleineren Abschnitt $B' < B$ von G einen ähnlichen Abschnitt $A' < A$ von F .‘ ...

H. ,Sind A und A' zwei Abschnitte einer wohlgeordneten Menge F , B und B' ihnen ähnliche Abschnitte einer wohlgeordneten Menge G , und ist $A' < A$, so ist $B' < B$.‘ ...

I. ,Ist ein Abschnitt B einer wohlgeordneten Menge G keinem Abschnitt einer wohlgeordneten Menge F ähnlich, so ist sowohl jeder Abschnitt $B' > B$ [$B' \supset B$] von G als auch G selbst weder einem Abschnitt von F noch F selbst ähnlich.‘ ...

[Es gibt keine Aussage J.]

K. ,Gibt es zu jedem Abschnitt A einer wohlgeordneten Menge F einen ihm ähnlichen Abschnitt B einer anderen wohlgeordneten Menge G , aber auch umgekehrt zu jedem Abschnitt B von G einen ihm ähnlichen Abschnitt A von F , so ist $F \approx G$ [$F \equiv G$].‘ ...

L. ,Gibt es zu jedem Abschnitt A einer wohlgeordneten Menge F einen ihm ähnlichen Abschnitt B einer anderen wohlgeordneten Menge G , ist hingegen mindestens ein Abschnitt von G vorhanden, zu dem es keinen ähnlichen Abschnitt von F gibt, so existiert ein bestimmter Abschnitt B_1 von G , so daß $B_1 \approx F$.‘ ...

M. ,Hat die wohlgeordnete Menge G mindestens einen Abschnitt, zu dem kein ähnlicher Abschnitt in der wohlgeordneten Menge F vorhanden ist, so muß jeder Abschnitt A von F einen ihm ähnlichen Abschnitt B in G haben.‘ ...

N. ,Sind F und G zwei beliebige wohlgeordnete Mengen, so sind entweder 1) F und G einander ähnlich, oder es gibt 2) einen bestimmten Abschnitt B_1 von G , welcher F ähnlich ist, oder es gibt 3) einen bestimmten Abschnitt A_1 von F , welcher G ähnlich ist; und jeder dieser drei Fälle schließt die Möglichkeit der beiden anderen aus.‘ ...

O. ,Ist eine Teilmenge F' einer wohlgeordneten Menge F keinem Abschnitt von F ähnlich, so ist sie F selbst ähnlich.‘ ... “

(Georg Cantor 1897, „Beiträge zur Begründung der transfiniten Mengenlehre“)

5. Der Wohlordnungssatz

Nachdem wir nun die möglichen Beziehungen von Wohlordnungen untereinander vollständig geklärt haben, ist die natürlichste Frage an dieser Stelle die nach der Existenz von Wohlordnungen auf beliebigen Mengen:

Kann jede Menge wohlgeordnet werden, d. h. existiert für jede Menge M eine Wohlordnung $<$ auf M ?

Wenn man etwa an die reellen Zahlen denkt, so ist keineswegs klar, wie eine Wohlordnung der reellen Zahlen aussehen soll. Cantor hat die Wohlordnenbarkeit jeder Menge zunächst als Denkgesetz postuliert, später hat er intuitive – und mit den Methoden der Nachfolgeneration streng zu rechtfertigende – Argumente für die Wohlordnenbarkeit jeder Menge gegeben.

Cantor (1883b): „Der Begriff der wohlgeordneten Menge weist sich als fundamental für die ganze Mannigfaltigkeitslehre aus. Daß es immer möglich ist, jede wohldefinierte Menge in die Form einer wohlgeordneten Menge zu bringen, auf dieses, wie mir scheint, grundlegende und folgenreiche, durch seine Allgemeingültigkeit besonders merkwürdige Denkgesetz werde ich in einer späteren Abhandlung zurückkommen.“

Der erste strenge Beweis der Wohlordnenbarkeit jeder beliebigen Menge gelang Ernst Zermelo im Jahre 1904. Er hat, gemessen am dreiseitigen Umfang der Veröffentlichung, für ein enormes Aufsehen in der mathematischen Welt gesorgt, wobei auch viele unsachliche Reaktionen nicht ausblieben.

Satz (*Wohlordnungssatz von Ernst Zermelo*)

Sei M eine Menge. Dann existiert eine Wohlordnung $<$ auf M .

Wir geben unten den originalen Beweis. Alle bekannten Beweise beruhen auf der Idee eines erschöpfenden Aufzählens aller Elemente der zugrunde liegenden Menge.

Es zeigt sich in der axiomatischen Entwicklung der Mengenlehre, daß der Wohlordnungssatz zu einem ausgezeichneten Axiom auf der Basis der übrigen Axiome äquivalent ist, nämlich dem Auswahlaxiom (dem Axiom, das Auswahlakte „ein ...“ ermöglicht). Insofern ist die Bezeichnung „Denkgesetz“ für den Wohlordnungssatz nicht unzutreffend.

Nun also zum Beweis des Wohlordnungssatzes!

Beweis

Für jede nichtleere Teilmenge A von M sei

$$\gamma(A) = \text{„ein } x \in A\text{“}.$$

[D.h. wir fixieren eine Funktion $\gamma : \mathcal{P}(M) - \{\emptyset\} \rightarrow M$ mit $\gamma(A) \in A$ für alle $A \in \text{dom}(\gamma) = \mathcal{P}(M) - \{\emptyset\}$.]

Eine Wohlordnung $\langle A, < \rangle$ heißt eine γ -Menge, falls gilt:

- (i) $A \subseteq M$,
- (ii) für alle $x \in A$ ist $x = \gamma(M - A_x)$,

wobei wieder $A_x = \{y \in A \mid y < x\}$ das durch x bestimmte Anfangsstück der Wohlordnung $\langle A, < \rangle$ bezeichnet. Insbesondere ist also $\gamma(M)$ das kleinste Element jeder γ -Menge $\langle A, < \rangle$ mit $A \neq \emptyset$. Weiter ist jedes Anfangsstück einer γ -Menge wieder eine γ -Menge.

Die Idee ist: Die γ -Mengen sind die Anfangsstücke einer bestimmten Wohlordnung von M , nämlich derjenigen Wohlordnung, die durch Abtragen von M gemäß γ entsteht. γ liefert uns an jeder Stelle des Abtragens ein Element aus dem Resthaufen, solange dieser noch nicht aufgebraucht ist.

Eine γ -Menge archiviert also bis zu einem bestimmten Zeitpunkt den Verlauf eines Prozesses, der ohne Willkür verläuft. Mit dieser Interpretation ist dann die folgende Aussage keine Überraschung:

- (+) Sind $\langle A, < \rangle$ und $\langle B, < \rangle$ zwei verschiedene γ -Mengen, so ist die eine ein Anfangsstück der anderen.

Beweis von (+)

Nach dem Vergleichbarkeitssatz für Wohlordnungen sei o.E.

$\langle A, < \rangle \equiv \langle B', < \rangle$, wobei $\langle B', < \rangle$ ein Anfangsstück von $\langle B, < \rangle$ oder gleich $\langle B, < \rangle$ sei.

Sei $\pi : A \rightarrow B'$ der zugehörige Ordnungsisomorphismus.

Es genügt zu zeigen: $\pi = \text{id}_A$.

Sei $X = \{x \in A \mid \pi(x) \neq x\}$.

Annahme $X \neq \emptyset$. Sei dann

$x = \text{„das kleinste Element von } X\text{“}$

und $z = \pi(x)$.

Nach minimaler Wahl von x ist dann $A_x = B'_z$.

Da $\langle A, < \rangle$ und $\langle B', < \rangle$ γ -Mengen sind, gilt also

$$x = \gamma(M - A_x) = \gamma(M - B'_z) = z, \text{ Widerspruch!}$$

Dies zeigt (+).

Sei $\Gamma = \{ \langle A, < \rangle \mid \langle A, < \rangle \text{ ist eine } \gamma\text{-Menge} \}$, und sei

$$\langle N, < \rangle = \bigcup \Gamma.$$

$\bigcup \Gamma$ ist eine Wohlordnung nach (+).

Weiter gilt:

(++) $\langle N, < \rangle$ ist eine γ -Menge.

Beweis von (++)

Offenbar gilt $N \subseteq M$.

Sei $x \in N$. Dann existiert eine γ -Menge $\langle A, < \rangle$ mit $x \in A$, also gilt $x = \gamma(M - A_x)$. Aber $A_x = N_x$ für alle $x \in A$, und damit $x = \gamma(M - N_x)$. Also ist $\langle N, < \rangle$ eine γ -Menge.

(+++)
Es gilt $N = M$.

Beweis von (+++)

Annahme $M - N \neq \emptyset$.

Sei $x = \gamma(M - N)$ und sei $\langle N', <' \rangle$ die Wohlordnung $\langle N, < \rangle$,
enderweitert um das Element x , d. h. $\langle N', <' \rangle = \langle N, < \rangle + \{x\}$.

Dann ist $\langle N', <' \rangle$ eine γ -Menge mit $x \in N'$, *im Widerspruch* zu $x \notin N$
und der Definition von N als Vereinigung der Träger aller γ -Mengen.

– Also ist $\langle N, < \rangle$ eine Wohlordnung auf M .

Zermelo hat 1908 einen weiteren Beweis des Wohlordnungssatzes gegeben, der die Theorie der Wohlordnungen nicht voraussetzt. Der Beweis des Vergleichbarkeitssatzes in 1.5 folgt der Struktur dieses zweiten Zermeloschen Beweises des Wohlordnungssatzes.

Die Stärke des Wohlordnungssatzes wollen wir noch mit einem neuen Beweis des Vergleichbarkeitssatzes für Mächtigkeiten demonstrieren:

Beweis des Vergleichbarkeitssatzes (1.5) mit Hilfe des Wohlordnungssatzes

Seien M, N Mengen. Wir zeigen $|M| \leq |N|$ oder $|N| \leq |M|$.

Nach dem Wohlordnungssatz existieren Wohlordnungen $\langle M, < \rangle$ und $\langle N, < \rangle$ von M und N . Weiter gilt nach dem Vergleichbarkeitssatz für Wohlordnungen:

$\langle M, < \rangle \triangleleft \langle N, < \rangle$ oder $\langle M, < \rangle \equiv \langle N, < \rangle$ oder $\langle N, < \rangle \triangleleft \langle M, < \rangle$.

Im ersten Fall existiert ein Ordnungsisomorphismus f von M auf ein Anfangsstück von N . Dann ist $f : M \rightarrow N$ injektiv, also gilt $|M| \leq |N|$.

Analog gilt im zweiten Fall $|M| = |N|$ und im dritten $|N| \leq |M|$.

– Dies zeigt die Behauptung.

Der Satz von Hartogs

Man kann umgekehrt den Vergleichbarkeitssatz von Mächtigkeiten verwenden, um den Wohlordnungssatz zu zeigen. Das Argument ist nicht trivial und braucht einen auch andernorts nützlichen Satz von Friedrich Hartogs (1874 – 1943) aus dem Jahre 1915. Hartogs ist heute eher bekannt für seine Arbeiten zur Funktionentheorie mehrerer komplexer Veränderlicher, aber seine Arbeit von

1915 gehört sicher zum Kulturerbe innerhalb der Mengenlehre. Sie bringt ein neues Argument für die Existenz langer Wohlordnungen. Schwächer als die Frage nach der Wohlordenbarkeit jeder Menge ist:

Kann es eine Menge geben, deren Mächtigkeit größer ist als die Mächtigkeit jeder wohlordenbaren Menge?

Wenn sich jede Menge wohlordnen läßt, so kann es eine solche Menge offenbar nicht geben. Die Frage läßt sich aber ohne den Wohlordnungssatz elementar verneinen, und zusammen mit dem Vergleichbarkeitssatz für Mächtigkeiten erhalten wir einen neuen Beweis für den Wohlordnungssatz.

Satz (Satz von Hartogs)

Sei M eine Menge. Dann existiert eine wohlordenbare Menge W derart, daß $\text{non}(|W| \leq |M|)$ gilt.

Wir schreiben $\text{non}(|W| \leq |M|)$ statt $|M| \leq |W|$, da wir den Vergleichbarkeitssatz für Mächtigkeiten nicht verwenden wollen. Ein W mit $\text{non}(|W| \leq |M|)$ ist genau das, was der Satz von Hartogs liefert.

Beweis

Wir setzen

$$H = \{ \langle A, < \rangle \mid \langle A, < \rangle \text{ ist eine Wohlordnung und } A \subseteq M \}.$$

$$W = H/\equiv.$$

W ist also die Menge der Äquivalenzklassen der Relation „gleichlang“ auf H .

Wir setzen für $\langle A, < \rangle/\equiv, \langle B, < \rangle/\equiv \in W$:

$$\langle A, < \rangle/\equiv < \langle B, < \rangle/\equiv \text{ falls } \langle A, < \rangle \text{ kürzer als } \langle B, < \rangle \text{ ist.}$$

Die Definition von $<$ auf $H/\equiv$ hängt nicht von der Wahl der Repräsentanten $\langle A, < \rangle$ und $\langle B, < \rangle$ ab. Weiter gilt:

(+) $<$ ist eine Wohlordnung auf $H/\equiv$.

Beweis von (+)

Linearität folgt aus dem Vergleichbarkeitssatz für Wohlordnungen.

Zur Wohlordnungseigenschaft:

Sei $J \subseteq H/\equiv$, $J \neq \emptyset$, und sei $\langle A, < \rangle/\equiv \in J$ beliebig. Wir setzen:

$$B = \{ x \in A \mid \langle A_x, < \rangle/\equiv \in J \}.$$

Ist $B = \emptyset$, so ist offenbar $\langle A, < \rangle/\equiv$ das $<$ -kleinste Element von J .

Andernfalls sei x das $<$ -kleinste Element von B in $\langle A, < \rangle$.

Dann ist $\langle A_x, < \rangle/\equiv \in J$ das $<$ -kleinste Element von J .

Weiter gilt offenbar für alle $\langle A, < \rangle/\equiv \in W$:

$$(++) \langle W_{\langle A, < \rangle/\equiv}, < \rangle = \langle \{ \langle A_x, < \rangle/\equiv \mid x \in A \}, < \rangle \equiv \langle A, < \rangle.$$

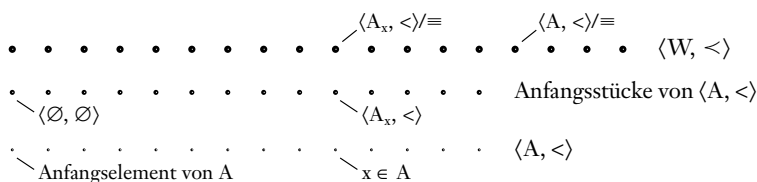
$[W_{\langle A, < \rangle/\equiv}$ ist das durch $\langle A, < \rangle/\equiv$ gegebene Anfangsstück von $\langle W, < \rangle$, also identisch mit $\{ Y \in W \mid Y < \langle A, < \rangle/\equiv \}$.]

Annahme, es gibt ein injektives $f: W \rightarrow M$.

Dann ist $\langle W, < \rangle \equiv \langle A, < \rangle$ für ein $\langle A, < \rangle \in H$, denn f induziert eine zu $\langle W, < \rangle$ gleichlange Wohlordnung auf $A = \text{rng}(f) \subseteq M$.

Nach (++) ist dann $\langle W, < \rangle \equiv \langle W_{\langle A, < \rangle / \equiv}, < \rangle$, also $\langle W, < \rangle \triangleleft \langle W, < \rangle$,

– *Widerspruch!* Also gilt $\text{non}(|W| \leq |M|)$.



Der Leser lasse sich durch die Flut von „ $\equiv$ “ nicht irritieren, der Beweis hat eine einfache Struktur. Sieht man jede Äquivalenzklasse von $H/\equiv$ als (dicken) Punkt an, so werden diese Punkte durch $<$ natürlicherweise wohlgeordnet. Jede Wohlordnung $\langle A, < \rangle$ mit $A \subseteq M$ repräsentiert über die – wieder in natürlicher Weise wohlgeordnete – Reihe $\langle A_x, < \rangle$, $x \in A$, ein Anfangsstück dieser Punkte, und repräsentiert selbst einen Punkt von $\langle W, < \rangle$. Da $\langle W, < \rangle$ wie jede Wohlordnung nicht gleichlang zu einem ihrer Anfangsstücke sein kann, kann kein $\langle A, < \rangle$ die Länge von $\langle W, < \rangle$ haben.

Übung

Sei M eine Menge. Dann existiert eine wohlordenbare Menge W derart, daß $\text{non}(|W| \leq^* |M|)$ gilt, d. h. es gibt kein surjektives $f: M \rightarrow W$.

[ohne „ein ...“: es gilt $|N| \leq^* |M|$ folgt $|N| \leq |\mathcal{P}(M)|$ für alle M, N .]

Definition (Hartogswohlordnung)

Sei M eine Menge. Dann heißt die im Beweis konstruierte Wohlordnung $\langle W, < \rangle$ die *Hartogswohlordnung von M* , in Zeichen $\mathcal{H}(M) = \langle W, < \rangle$.

Mit dem Satz von Hartogs folgt nun leicht:

Beweis des Wohlordnungssatzes mit Hilfe des Vergleichbarkeitssatzes (1.5)

Sei M eine Menge, und sei $\mathcal{H}(M) = \langle W, < \rangle$ die Hartogswohlordnung von M .

Dann gilt $\text{non}(|W| \leq |M|)$. Nach dem Vergleichbarkeitssatz für Mächtigkeiten gilt also $|M| \leq |W|$. Sei also $f: M \rightarrow W$ injektiv.

– Dann ist $< = \{ (x, y) \in M \times M \mid f(x) < f(y) \}$ eine Wohlordnung auf M .

Hartogs (1915): „Im folgenden gebe ich für den Satz, daß jede Menge wohlgeordnet werden kann, einen Beweis, der sich von den beiden Zermeloschen Beweisen [1904, 1908] dadurch unterscheidet, daß das sogenannte Auswahlprinzip [Definitionen der Form ‚ein ...‘] bei ihm nicht zur Anwendung kommt, dafür jedoch eine Prämisse anderer Art benutzt wird, nämlich die Annahme der ‚Vergleichbarkeit der Mengen‘...

Wird von den drei genannten Prinzipien [Auswahlprinzip, Vergleichbarkeit, Wohlordenbarkeit] keines vorausgesetzt, so liefern die folgenden Betrachtungen immer noch einen Nachweis für den Satz, daß es keine Menge geben kann, deren Mächtigkeit größer ist als die Mächtigkeit jeder beliebigen wohlgeordneten Menge**)

**) Dieser Satz ist meines Wissens ohne Anwendung des Auswahlprinzips bisher noch nicht streng bewiesen worden.“

In diesem Zusammenhang ist eine Fußnote von Arthur Schoenflies (1853 – 1928) nicht ohne Witz, und Hartogs bezieht sich möglicherweise auf diese Stelle:

Schoenflies (1908, S. 32f.): „Es liegt nahe, Mengen, deren Mächtigkeit jedes Aleph übertrifft [deren Mächtigkeit die jeder wohlordenbaren Menge übertrifft], als widerspruchsvoll zu betrachten. Aber dies hängt andererseits so eng mit den nicht hinlänglich geklärten Problemen der Wohlordnung und der Vergleichbarkeit zusammen, daß ich vorziehe, mich dahin auszusprechen, es liege hier noch eine Lücke der Erkenntnis vor. Immerhin scheint es mir zweckmäßig jede Berufung auf die Existenz oder Nichtexistenz derartiger Mengen zu vermeiden.“

1908 war das Problem der Wohlordnung und der Vergleichbarkeit eigentlich schon geklärt, aber vielerorts bestand immer noch Skepsis gegenüber Zermelos Beweis des Wohlordnungssatzes. Schoenflies diagnostiziert eine Lücke der Erkenntnis, vermutet sie aber am falschen Ort: Hartogs zeigt, daß keine Menge existiert, deren Mächtigkeit hinter dem Reich der wohlordenbaren Mengen liegt, und sein Argument der Verschmelzung aller auf Teilmengen einer Menge existierenden Wohlordnungen zu einer neuen Wohlordnung hat mit dem Vergleichbarkeitssatz nichts zu tun, und es ist frei von abstrakten Auswahlakten.

Eine weitere Anwendung des Satzes von Hartogs ist ein elementarer Beweis des Wohlordnungssatzes mit Hilfe der Blackbox Multiplikationssatz.

Beweis des Wohlordnungssatzes mit Hilfe des Multiplikationssatzes

Sei M eine beliebige Menge, und sei $\langle W, < \rangle$ eine Wohlordnung mit $\text{non}(|W| \leq |M|)$. Ohne Einschränkung ist W unendlich. Dann gilt

$$|M \times W| \leq |M|^2 \cup M \times W \cup W \times M \cup W^2 = |(M \cup W)^2| = |M \cup W|,$$

wobei im letzten Schritt der Multiplikationssatz verwendet wird.

Nach dem Satz von Bernstein (1.5) gilt also $|M| \leq |W|$ oder $|W| \leq |M|$.

$|W| \leq |M|$ ist nach Wahl von W unmöglich, also gilt $|M| \leq |W|$.

- Wie oben induziert eine Injektion $f: M \rightarrow W$ eine Wohlordnung auf M .

Der Beweis ist elementar in dem Sinne, daß Auswahlakte nicht verwendet werden müssen: Der Leser, der den Beweis des Satzes von Bernstein in 1.5 betrachtet, wird sehen, daß die dortige Verwendung von „ein ...“ eine Auswahl innerhalb der Menge N betrifft. Hier ist $N = W$ wohlordenbar, sodaß „ein ...“ durch „das <-kleinste ...“ ersetzt werden kann.

Daß die Kombination des Satzes von Bernstein mit dem Satz von Hartogs den Wohlordnungssatz elementar aus dem Multiplikationssatz liefert, hat Tarski bemerkt [Tarski 1924]. Der Trick mit der binomischen Gleichung findet sich ebenfalls schon in der Dissertation von Bernstein 1901 [vgl. Bernstein 1905, S. 132].

Der Beweis zeigt eine lokale Version des Satzes: Gilt $|(M \cup W)^2| = |M \cup W|$ für eine einzige Wohlordnung W , die mindestens so lang ist wie die Hartogswohlordnung $\mathcal{H}(M)$, so ist M wohlordenbar.

In der axiomatischen Mengenlehre lesen sich unsere „mit Hilfe von ...“-Resultate dann insgesamt wie folgt: Auswahlaxiom, Wohlordnungssatz, Vergleichbarkeitssatz und Multiplikationssatz sind über den restlichen Axiomen äquivalent (vgl. hierzu Abschnitt 3).

Wir werden den Satz von Hartogs im übernächsten Kapitel noch einmal verwenden.

Maximalprinzipien

In den Beweisen des Vergleichbarkeitssatzes (1.5) und des Wohlordnungssatzes haben wir die Existenz der gesuchten Objekte durch den Nachweis gezeigt, daß bestmögliche Approximationen an diese Objekte existieren. Die hier verwendete Argumentation läßt sich in abstrakten Prinzipien bündeln. Im Satz über die Existenz von Zielen in Zermelosystemen haben wir schon ein Beispiel kennengelernt, zwei weitere Varianten werden wir nun besprechen, den Satz von Zermelo-Zorn und das Maximalitätsprinzip von Hausdorff.

Der Satz von Zermelo-Zorn wird unter dem Pseudonym „Zornsches Lemma“ an verschiedenen Stellen in der mathematischen Praxis gerne eingesetzt (etwa im Beweis der Existenz von Basen in beliebigen Vektorräumen, im Beweis des Produktsatzes von Tychonov in der Topologie, im Existenzbeweis des algebraischen Abschlusses eines Körpers in der Algebra, im Beweis des Satzes von Hahn-Banach in der Funktionalanalysis). Zur Formulierung brauchen wir den Begriff der partiellen Ordnung, den wir bei der Einführung der linearen Ordnungen schon kurz erwähnt haben:

Definition (*partielle Ordnung*)

Sei M eine Menge, und sei $< \subseteq M \times M$.

$\langle M, < \rangle$ heißt eine *partielle Ordnung* auf M , falls $<$ irreflexiv und transitiv auf M ist.

Eingeführt wurde der Begriff der partiellen Ordnung von Hausdorff 1914. Er gehört heute zum Grundbestand des mathematischen Wortschatzes.

Hausdorff (1914): „Nehmen wir an, zwischen je zwei verschiedenen Elementen a, b einer Menge A bestehe jetzt nicht mehr, wie bei den [linear] geordneten Mengen, eine und nur eine von zwei Beziehungen ($a < b, a > b$), sondern eine und nur eine von drei Beziehungen

$$a < b, a > b, a \parallel b,$$

die wir lesen wollen: a vor b , a nach b , a unvergleichbar mit b . Von den beiden ersten setzen wir dieselben Eigenschaften wie im Falle geordneter Mengen voraus, was für die dritte Beziehung notwendig ihre Symmetrie zur Folge hat [d. h. $a \parallel b$ folgt $b \parallel a$] ...

Eine solche Menge heißt eine teilweise geordnete Menge; die geordneten Mengen sind Spezialfälle der teilweise geordneten, nämlich wenn Paare unvergleichbarer Elemente nicht existieren ...“

Man nennt Elemente a, b einer partiellen Ordnung mit $a \parallel b$, d. h. $\text{non}(a < b)$ und $\text{non}(b < a)$ auch *inkompatibel*. Umgangssprachlich ist ein Analogon zur mathematischen partiellen Ordnung gut bekannt: „Man kann Äpfel nicht mit Birnen vergleichen“, aber ein Apfel kann größer sein als ein anderer. Die Ergebnisse von olympischen Spielen mit ihren Dutzenden von Sportarten bilden ein Beispiel für eine sinnvolle partiell geordnete Struktur. Jeder Baum gibt ein schönes.

Wir verwenden die Schreibweisen für lineare Ordnungen auch für partielle Ordnungen. So meint $A \leq x$ etwa „ $a \leq x$ für alle $a \in A$ “, und x ist dann eine *obere Schranke* von $A \subseteq M$ in der partiellen Ordnung $\langle M, < \rangle$.

Partielle Ordnungen haben eine netzartige Struktur. Die beiden wichtigsten Beispiele sind die Inklusion $\subset$ auf einer beliebigen Menge M , d. h. $a < b$ falls $a \subset b$ für $a, b \in M$, sowie die umgekehrte Inklusion $\supset$ auf einer Menge M , d. h. $a < b$ falls $a \supset b$ für $a, b \in M$. (Häufig ist $M = \mathcal{P}(N)$ für eine Menge N .)

Weiter definieren wir:

Definition (*Ziele und Ketten in einer partiellen Ordnung*)

Sei $\langle M, < \rangle$ eine partielle Ordnung.

- (a) $x \in M$ heißt ein *Ziel* oder ein *maximales Element* von $\langle M, < \rangle$, falls kein $y \in M$ existiert mit $x < y$.
- (b) $K \subseteq M$ heißt eine *Kette* oder *linear geordnete Teilmenge* von $\langle M, < \rangle$, falls $\langle K, <|_K \rangle$ eine lineare Ordnung ist.

Ist $\langle M, < \rangle = \langle \mathcal{P}(\mathbb{N}), \subset \rangle$, so ist $\mathbb{N}$ das einzige Ziel von $\langle M, < \rangle$. Die partiellen Ordnungen $\langle \mathbb{N}, < \rangle$ und $\langle \{x \subseteq \mathbb{N} \mid x \text{ endlich} \}, \subset \rangle$ haben keine Ziele. Im allgemeinen sind Ziele in einer partiellen Ordnung in einer Vielzahl vorhanden:

Übung

Bestimmen Sie die Ziele der folgenden partiellen Ordnungen:

- (i) $\langle \mathcal{P}(\mathbb{N}) - \{\mathbb{N}\}, \subset \rangle$,
- (ii) $\langle (\mathcal{P}(\mathbb{N}))^2, < \rangle$, wobei $<$ für alle $x_1, x_2, y_1, y_2 \subseteq \mathbb{N}$ definiert ist durch:
 $(x_1, x_2) < (y_1, y_2)$, falls $x_1 \subset y_1$ und $x_2 \subset y_2$.

Das „Zornsche Lemma“ garantiert die Existenz von Zielen für eine Vielzahl von partiellen Ordnungen. Der Beweis ergibt sich leicht aus dem Satz über die Existenz von Zielen in Zermelosystemen.

Satz (*Satz von Zermelo-Zorn, „Zornsches Lemma“*)

Sei $\langle M, < \rangle$ eine partielle Ordnung. Für alle Ketten K in M existiere eine obere Schranke in M , d. h. ein $x \in M$ mit $K \leq x$.

Dann existiert für alle $x_0 \in M$ ein Ziel x von M mit $x_0 \leq x$.

Beweis

Für $x \in M$ sei $V_x = \{y \in M \mid y \leq x\}$. Dann gilt für alle $x, y \in M$:

$x < y$ gdw $V_x \subset V_y$.

Sei $\mathcal{V} = \{V_x \mid x \in M\}$. $\mathcal{V}$ ist i. a. kein Zermelosystem, aber wir können $\mathcal{V}$ zu einem Zermelosystem erweitern. Wir setzen hierzu:

$\mathcal{L} = \{\bigcup K \mid K \subseteq \mathcal{V} \text{ ist eine Kette}\}.$

Dann ist $\mathcal{L} \supseteq \mathcal{V}$ (für jedes $V_x \in \mathcal{V}$ ist $\{V_x\}$ eine Kette, und $\bigcup \{V_x\} = V_x$).

Weiter ist $\mathcal{L}$ ein Zermelosystem (!). Sei also $V \supseteq V_{x_0}$ ein Ziel von $\mathcal{L}$.

Dann gilt:

(+) $V \in \mathcal{V}$, d.h. es gibt ein $x \in M$ mit $V = V_x$.

Denn sei $K \subseteq \mathcal{V}$ eine Kette mit $V = \bigcup K$.

Sei $K' = \{y \in M \mid V_y \in K\}$. Dann ist K' eine Kette in $\langle M, < \rangle$.

Nach Voraussetzung an $\langle M, < \rangle$ existiert ein x mit $K' \leq x$.

Dann ist $W \subseteq V_x$ für alle $W \in K$. Also $V = \bigcup K \subseteq V_x$.

Aber V ist ein Ziel von $\mathcal{L}$, also $V = V_x$. Dies zeigt (+).

- Dann ist x aber ein Ziel von $\langle M, < \rangle$, denn gäbe es ein $y \in M$ mit $x < y$,
 – so wäre $V_x \subset V_y$, also wäre V_x kein Ziel von $\mathcal{L}$. Zudem gilt $x_0 \leq x$.

Die leere Menge gilt als linear geordnete Teilmenge (also als Kette), hat also unter der Voraussetzung an die partielle Ordnung eine obere Schranke $x \in M$. Damit ist immer $M \neq \emptyset$, falls die obere Schranken-Bedingung für Ketten erfüllt ist.

Umgekehrt folgt aus dem Zornschen Lemma der Satz über die Existenz von Zielen in Zermelosystemen: Ist $\mathcal{L}$ ein Zermelosystem, so erfüllt $\langle \mathcal{L}, \subset \rangle$ die Voraussetzung des Zornschen Lemmas, und jedes Ziel von $\langle \mathcal{L}, \subset \rangle$ ist ein Ziel von $\mathcal{L}$.

Ein Beispiel für eine partielle Ordnung, deren Träger kein Zermelosystem bildet, für die aber dennoch die Schrankenbedingung gilt, ist etwa $\langle \mathcal{F}, \subset \rangle$ mit $\mathcal{F} = \{X \subseteq \mathbb{R} \mid X \text{ ist abgeschlossen}\}$.

Historisch ist das „Zornsche Lemma“ für Zermelosysteme formuliert worden. Es findet sich in einer Arbeit von Max Zorn aus dem Jahre 1935. Das Prinzip wird dort ohne Beweis angegeben, was zu einer Art Tradition in der Mathematik geworden ist. Es hat reinen Werkzeugcharakter, daher auch die Bezeichnung „Lemma“. Es ermöglicht, gewisse Beweise auch ohne Kenntnis der Wohlordnungstheorie führen zu können. Der Beweis des Satzes von Zermelo-Zorn ruht aber vollkommen auf der etwa dreißig Jahre früher entwickelten Technik von Zermelo – in den Varianten von 1904 mit Wohlordnungen wie im Beweis des Wohlordnungssatzes oben, und von 1908 ohne Wohlordnungen wie in 1.5 –, sodaß die heute weitverbreitete Namensgebung nicht ganz glücklich erscheint. Man sagt ja auch Satz von Cantor-Bernstein, ohnehin schon Dedekind verschluckend. Warum also nicht Satz von Zermelo-Zorn? Den meisten Mathematikern ist heute Zermelo bestenfalls als „Mister Z“ der Zermelo-Fraenkel-Axiomatik bekannt. Es geht nicht so sehr um Ruhm und Ehre. Ungenaue Zuordnungen von Mensch und Begriff vernebeln die Geschichte.

Wir geben noch einen zweiten Beweis des Satzes von Zermelo-Zorn, der fast zeilenweise dem Zermeloschen Beweis des Wohlordnungssatzes folgt.

Zweiter Beweis des Satzes von Zermelo-Zorn

Sei $\mathcal{P}(M)^* = \{A \subseteq M \mid \text{es existiert ein } x \in M \text{ mit } A < x\}$.

Für $A \in \mathcal{P}(M)^*$ sei $\sigma(A) = \text{„ein } x \in M \text{ mit } A < x\text{“}$.

Damit ist $\sigma : \mathcal{P}(M)^* \rightarrow M$ eine Funktion, die uns – im Falle der Existenz – echte obere Schranken für Teilmengen von M liefert.

Wir nehmen weiter $\sigma(\emptyset) = x_0$ an.

Eine Wohlordnung $\langle A, < \rangle$ heißt eine σ -Menge, falls gilt:

(i) $A \subseteq M$,

(ii) für alle $x \in A$ ist $x = \sigma(A_x)$, wobei $A_x = \{y \in A \mid y < x\}$ für $x \in A$.

Wie im Beweis des Wohlordnungssatzes zeigt man, daß die Vereinigung $\langle N, < \rangle$ aller σ -Mengen eine σ -Menge ist.

Sei x eine obere Schranke für N , also $N \leq x$. Eine solche Schranke existiert nach Voraussetzung an $\langle N, < \rangle$, denn N ist eine linear geordnete (sogar wohlgeordnete) Teilmenge von M .

Andererseits gilt $N \notin \text{dom}(\sigma)$. Denn *andernfalls* ist $N < \sigma(N)$, und dann ist

$$\langle N', <' \rangle = \langle N, < \rangle + \{ \sigma(N) \}$$

eine σ -Menge. Folglich $N' \subseteq N$ nach Definition von N .

Also $\sigma(N) \in N$, *Widerspruch!*

Wegen $N \notin \text{dom}(\sigma)$ existiert keine echte obere Schranke für N .

Andererseits gilt $N \leq x$ nach Wahl von x .

- Also ist x ein maximales Element von M . Weiter gilt $x_0 \leq x$.

Übung

Beweisen Sie den Wohlordnungssatz mit Hilfe des Satzes von Zermelo-Zorn (oder mit Hilfe des Satzes über die Existenz von Zielen in Zermelosystemen).

Das Hausdorffsche Maximalitätsprinzip

Hausdorff hat bereits 1914 ein sehr hochwertiges Maximalitätsprinzip betrachtet [Hausdorff 1914, S. 140 f.]. Das Prinzip handelt nicht von Zielen in partiellen Ordnungen, sondern von maximalen Ketten in ihnen. Es folgt durch Konstruktion einer zweiten partiellen Ordnung, bestehend aus den Ketten der ersten, leicht aus dem Satz von Zermelo-Zorn, und umgekehrt ergibt sich der Satz von Zermelo-Zorn leicht aus diesem Kettenprinzip.

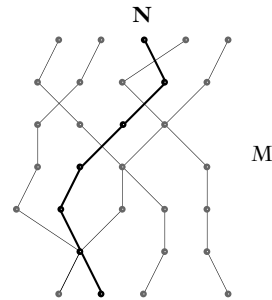
Satz (Hausdorffsches Maximalitätsprinzip)

Sei $\langle M, < \rangle$ eine partielle Ordnung.

Dann existiert eine maximale Kette in M ,

d. h. es gibt ein $N \subseteq M$ mit der Eigenschaft:

- (i) $\langle N, <|_N \rangle$ ist linear geordnet.
- (ii) Für kein $N' \subseteq M$ mit $N \subset N'$ ist $\langle N', <|_{N'} \rangle$ linear geordnet.



Übung

- (A) Beweisen Sie das Hausdorffsche Prinzip:
 - (a) analog zum Beweis des Wohlordnungssatzes,
 - (b) mit Hilfe des Satzes über Zermelosysteme.
- (B) Beweisen Sie den Satz von Zermelo-Zorn mit Hilfe des Hausdorffschen Prinzips.

Hausdorff (1914): „Eine teilweise [partiell] geordnete Menge A hat (vollständig) geordnete [linear geordnete] Teilmengen, z.B. mindestens die aus einem Element bestehenden. Eine geordnete Teilmenge, die in keiner anderen geordneten Teilmenge als echte Teilmenge enthalten ist, also nicht durch Hinzunahme anderer Elemente zu einer geordneten Teilmenge erweitert werden kann, nennen wir eine größte [maximale] geordnete Teilmenge. Die Existenz solcher werden wir zu beweisen haben.“

Schließlich erwähnen wir noch ein weiteres Maximalitätsprinzip, bekannt als *Teichmüller-Tukey Lemma* [Teichmüller 1939], [Tukey 1940].

Satz (*Teichmüller-Tukey-Lemma*)

Sei T eine nichtleere Menge. Für alle x gelte:

(+) $x \in T$ gdw für jedes endliche $y \subseteq x$ gilt $y \in T$.

Dann hat T ein Ziel, d.h. es gibt ein $x \in T$ mit: $\text{non}(x \subset y)$ für alle $y \in T$.

Übung

- (a) Beweisen Sie das Lemma von Teichmüller-Tukey mit Hilfe des Satzes von Zermelo-Zorn.
- (b) Beweisen Sie den Satz von Zermelo-Zorn oder das Hausdorffprinzip mit Hilfe des Lemmas von Teichmüller-Tukey.

Genug der Maximalprinzipien. Entscheidend für die Mengenlehre ist der Wohlordnungssatz, und das zweite fundamentale mengentheoretische Resultat dieses Kapitels ist der Satz von Hartogs. Die Maximalprinzipien finden in der Mengenlehre etwas weniger Anwendung als anderswo, da hier Wohlordnungen zur Verfügung stehen. Am klarsten und elegantesten werden die Beweise des Wohlordnungs- und Vergleichbarkeitssatzes und der Maximalprinzipien bei Verwendung der Ordinalzahlen und der transfiniten Rekursion, die wir in den beiden folgenden Kapiteln besprechen.

Ernst Zermelos Wohlordnungssatz

„Beweis, daß jede Menge wohlgeordnet werden kann.

(Aus einem an Herrn Hilbert gerichteten Briefe.)

Von E. ZERMELO in Göttingen

... Der betreffende Beweis ist aus Unterhaltungen entstanden, die ich in der vorigen Woche mit Herrn Erhard Schmidt geführt habe, und ist folgender.

1) Es sei M eine beliebige Menge ...

2) *Jeder [nichtleeren] Teilmenge M' [von M] denke man sich ein beliebiges Element m_1' zugeordnet, das in M' selbst vorkommt und das ‚ausgezeichnete‘ Element von M' genannt werden möge.* So entsteht eine ‚Belegung‘ γ der Menge M [$\mathcal{P}(M) - \{\emptyset\}$] mit Elementen der Menge M von besonderer Art ... Im folgenden wird nun eine beliebige Belegung γ zugrunde gelegt und aus ihr eine bestimmte Wohlordnung der Elemente von M abgeleitet.

3) *Definition.* Als ‚ γ -Menge‘ werde bezeichnet jede wohlgeordnete Menge $M_\gamma \subseteq M$, welche folgende Beschaffenheit besitzt: ist a ein beliebiges Element von M_γ und A der ‚zugehörige‘ Abschnitt, der aus den vorangehenden Elementen $x < a$ von M_γ besteht, so ist a immer das ‚ausgezeichnete Element‘ von $M - A$.

4) *Es gibt γ -Mengen innerhalb M .* So ist z. B. m_1 , das ausgezeichnete Element von $M' = M$, selbst eine γ -Menge ...

5) Sind M_γ' und M_γ'' irgend zwei verschiedene γ -Mengen (die aber zu derselben ein für allemal gewählten Belegung gehören!), so ist immer eine von beiden identisch mit einem Abschnitt der anderen ...

6) *Folgerungen.* Haben zwei γ -Mengen ein Element a gemeinsam, so haben sie auch den Abschnitt A der vorangehenden Elemente gemein. Haben sie zwei Elemente a, b gemein, so ist in beiden Mengen entweder $a < b$ oder $b < a$.

7) Bezeichnet man als ‚ γ -Element‘ jedes Element von M , das in irgendeiner γ -Menge vorkommt, so gilt der Satz: *Die Gesamtheit L_γ aller γ -Elemente läßt sich so ordnen, daß sie selbst eine γ -Menge darstellt, und umfaßt alle Elemente der ursprünglichen Menge M .* Die letztere ist damit selbst wohlgeordnet ...

Somit entspricht jeder Belegung γ eine ganz bestimmte Wohlordnung der Menge M ...

Der vorliegende Beweis beruht auf der Voraussetzung, daß Belegungen γ überhaupt existieren, also auf dem Prinzip, daß es auch für eine unendliche Gesamtheit von Mengen immer Zuordnungen gibt, bei denen jeder Menge eines ihrer Elemente entspricht ...

Die Idee, unter Berufung auf dieses Prinzip eine beliebige Belegung γ der Wohlordnung zugrunde zu legen, verdanke ich Herrn Erhard Schmidt; meine Durchführung des Beweises beruht dann auf der Verschmelzung der verschiedenen möglichen ‚ γ -Mengen‘, d. h. der durch das Ordnungsprinzip sich ergebenden wohlgeordneten Abschnitte.

Münden i. Hann., den 24. September 1904.“

(Ernst Zermelo 1904, „Beweis, daß jede Menge wohlgeordnet werden kann“)

6. Ordinalzahlen

Wir haben gezeigt, daß je zwei Wohlordnungen ihrer Länge nach vergleichbar sind. Die Idee der Cantorschen Ordinalzahlen ist nun, Wohlordnungen nur auf ihre Länge hin zu betrachten und dabei von der Natur des Trägers der Ordnung ganz abzusehen. Cantor definiert Ordinalzahlen wie folgt:

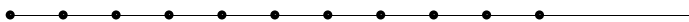
Cantor (1897):

„§ 14. Die Ordnungszahlen wohlgeordneter Mengen.

Nach § 7 hat jede einfach [linear] geordnete Menge M einen bestimmten *Ordnungstypus* $\bar{M}$; es ist dies der Allgemeinbegriff, welcher sich aus M ergibt, wenn unter Festhaltung der Rangordnung ihrer Elemente von der Beschaffenheit der letzteren abstrahiert wird, so daß aus ihnen lauter Einsen werden, die in einem bestimmten Rangverhältnis zu einander stehen. *Allen einander ähnlichen Mengen, und nur solchen, kommt ein und derselbe Ordnungstypus zu.*

Den Ordnungstypus einer wohlgeordneten Menge F nennen wir die ihr zukommende „*Ordnungszahl*“.

Das Einsetzen der Einsen veranschaulicht sehr gut den Prozeß der Abstraktion von den Trägerelementen, und entspricht recht genau den Skizzen von Wohlordnungen, die wir auf das Papier malen: Die Elemente der Wohlordnungen solcher Diagramme sind meistens ununterscheidbare Punkte:



Wer aber an die heute übliche, durch die Sprache der Mengenlehre mitbegründete Genauigkeit der Definitionen mathematischer Objekte gewöhnt ist, für den haftet wahrscheinlich einer Begriffsbildung dieser Art ein etwas fader Beigeschmack an.

Wir haben das gleiche Problem wie früher mit dem Begriff der Mächtigkeit vor uns: Es ist einfach, „gleichmächtig“ und „kleiner“ bzw. „gleichlang“ und „kürzer“ relational für zwei Mengen bzw. Wohlordnungen zu definieren. Von der „Mächtigkeit“ einer Menge oder der „Länge“ einer Wohlordnung schlechthin zu sprechen, ist schwieriger.

Die „universalen“ rein extensionalen Definitionen:

$\text{Mächtigkeit}(M) = \{ N \mid N \text{ und } M \text{ sind gleichmächtig} \},$

$\text{Länge}(\langle M, < \rangle) = \{ \langle N, < \rangle \mid \langle N, < \rangle \text{ und } \langle M, < \rangle \text{ sind gleichlang} \},$

sind ungeeignet, da die Zusammenfassungen für $M \neq 0$ zu groß sind, um noch Mengen sein zu können; es sind echte Klassen im Sinne der Diskussion in 1.13.

In der axiomatischen Mengenlehre kann eine Modifikation dieser Definitionen durchgeführt werden: Man kann die echten Klassen Mächtigkeit(M) und Länge($\langle M, < \rangle$) mit Hilfe eines Rangbegriffs uniform zu Mengen zurechtstutzen.

Die Cantorsche Definition zeigt, daß Cantor keineswegs mit „universale“ oder „Allgemeinbegriff“ obige Klassendefinition im Auge hatte. Durch Abstraktion entsteht ein gewisses Objekt aus Einsen, das nicht mehr der rein extensionalen Mengenwelt angehört. Wesentlich ist, daß die Abstraktion angewendet auf zwei Wohlordnungen genau dann das gleiche Objekt erzeugt, wenn die beiden Wohlordnungen gleichlang sind. Diese Eigenschaft ist es, die zählt, und die für mathematische Argumente relevant ist. (Der Leser vergleiche die Definition des geordneten Paares in 1.3 und ihre fundamentale Eigenschaft (#).)

Hausdorff sieht von Cantors Abstraktion ab. Er nimmt wie schon für Kardinalzahlen einen „formalen Standpunkt“ ein und sieht Ordinalzahlen als Zeichen an. Allgemein definiert er für lineare Ordnungen:

Hausdorff (1914): „Das den ähnlichen Mengen [ordnungsisomorphen linearen Ordnungen] Gemeinsame bezeichnen wir als Ordnungstypus, wie wir das den äquivalenten [gleichmächtigen] Mengen Gemeinsame als Mächtigkeit bezeichneten. Wir ordnen nämlich jeder [linear geordneten] Menge A ein Zeichen α zu, derart, daß ähnlichen Mengen und nur solchen dasselbe Zeichen entspricht, daß also mit $A \simeq B$ [$\langle A, < \rangle \equiv \langle B, < \rangle$] zugleich $\alpha = \beta$ und umgekehrt mit $\alpha = \beta$ zugleich $A \simeq B$ ist. Dieses Zeichen α heißt der Ordnungstypus (oder Typus) der Menge A .“

Ordinalzahlen definiert Hausdorff als die Ordnungstypen von Wohlordnungen.

So unterschiedlich die philosophischen und pragmatischen Anteile der Definitionen von Cantor und Hausdorff sein mögen, so identisch erweisen sie sich für die mathematische Praxis. Sie sind streng genommen keine Figuren auf unserer Bühne, aber man kann hervorragend mit ihnen kommunizieren. Wie im Fall der Kardinalzahlen sind Ordinalzahlen prinzipiell aus Resultaten eliminierbar – was nie gemacht wird, aber garantiert, daß wir korrekte Ergebnisse auch dann erzielen, wenn wir eine naive Definition zugrundelegen.

Wir arbeiten zunächst mit einer Definition nach Cantor, oder, mathematisch gleichwertig, nach Hausdorff. Wir gelangen dabei auf ganz natürlichem Weg zur heute üblichen Definition einer Ordinalzahl nach von Neumann und Zermelo, der keinerlei Vagheiten mehr anhaften, und die dann auch einen Schleifstein für den noch unscharfen Kardinalzahlbegriff im Gepäck hat. Die moderne Definition hat ohne Vorbereitung einen ad-hoc-Charakter, nach einer experimentellen Phase mit Cantor-Hausdorffschen Ordinalzahlen erscheint sie zwingend.

Die ideelle Zuordnung von Ordnungstypen werden wir aber bei der Diskussion allgemeiner linearer Ordnungen erneut verwenden, wo sie wieder, wenn man so will, einfach als eine bequeme Sprechweise aufgefaßt werden kann. Hier bleibt auch der axiomatischen Mengenlehre nur die Idee des Zurückschneidens durch Rangbetrachtungen. Die modernen Ordinalzahlen sind kanonische und definierbare Vertreter für alle Isomorphieklassen bestehend aus Wohlordnungen. Derartige Vertreter konnten allgemein für lineare Ordnungen nicht gefunden werden.

Die Definition nach Cantor und Hausdorff

Für das Folgende denken wir uns also jeder Wohlordnung $\langle M, < \rangle$ einen *Ordnungstyp* o. t. $\langle \langle M, < \rangle \rangle$ [o. t. für engl. order type] in einer solchen Weise zugeordnet, daß für je zwei Wohlordnungen $\langle M, < \rangle$ und $\langle N, < \rangle$ gilt:

o. t. $\langle \langle M, < \rangle \rangle =$ o. t. $\langle \langle N, < \rangle \rangle$ gdw $\langle M, < \rangle$ und $\langle N, < \rangle$ sind gleichlang.

Mit dieser Hypothek definieren wir:

Definition (Ordinalzahl)

Die Ordnungstypen von Wohlordnungen heißen *Ordinalzahlen*.

Neben „Ordinalzahl“ ist auch „Ordnungszahl“ gebräuchlich. Wir verwenden vorwiegend kleine griechische Buchstaben für Ordinalzahlen. Ist α eine Ordinalzahl und $\langle M, < \rangle$ eine Wohlordnung des Typs α , so sagen wir auch, daß $\langle M, < \rangle$ den Typ α repräsentiert.

Bei der Zuordnung von Ordnungstypen zu Wohlordnungen haben wir große Freiheiten. Es ist nun – wie schon für die Kardinalzahlen – sehr bequem, für endliche Wohlordnungen die natürlichen Zahlen als Ordnungstypen festzusetzen. Für $n \in \mathbb{N}$ sei wieder $\bar{n} = \{0, 1, \dots, n-1\}$, und es sei $<$ die übliche Ordnung auf $\bar{n}$. Dann gilt per Vereinbarung für alle $n \in \mathbb{N}$:

o. t. $\langle \langle \bar{n}, < \rangle \rangle = n$.

Die natürlichen Zahlen werden damit zu einem Anfangsstück der Ordinalzahlen, denn es gibt keine Wohlordnungen $\langle M, < \rangle$ mit $\langle \bar{n}, < \rangle \triangleleft \langle M, < \rangle \triangleleft \langle \bar{m}, < \rangle$ für $n \in \mathbb{N}$ und $m = n + 1$. Die triviale Wohlordnung $\langle \emptyset, \emptyset \rangle$ hat den Ordnungstyp 0. Für jedes Objekt x ist $\langle \{x\}, \emptyset \rangle$ eine Wohlordnung vom Typ 1, für je zwei verschiedene Objekte x, y ist $\langle \{x, y\}, \{(x, y)\} \rangle$ eine Wohlordnung vom Typ 2. Für jede endliche Wohlordnung $\langle M, < \rangle$ sind der Ordnungstyp und die Kardinalität der Trägermenge M identisch.

Weiter verwenden wir das nun schon vertraute Symbol ω für den durch die natürlichen Zahlen repräsentierten Ordnungstyp. Also:

$\omega =$ o. t. $\langle \langle \mathbb{N}, < \rangle \rangle$,

mit der üblichen Ordnung $<$ auf $\mathbb{N}$. Es ist natürlich, konkret $\mathbb{N}$ selbst als das durch ω bezeichnete Objekt zu wählen, und dann gilt $\omega = \mathbb{N} = \aleph_0$. Wie erwartet gilt:

Übung

ω ist die kleinste Ordinalzahl, die größer ist als alle $n \in \mathbb{N}$, d. h.:

Ist $\langle M, < \rangle$ eine Wohlordnung mit $\langle \bar{n}, < \rangle \triangleleft \langle M, < \rangle$ für alle $n \in \mathbb{N}$,

so gilt $\langle \mathbb{N}, < \rangle \equiv \langle M, < \rangle$ oder $\langle \mathbb{N}, < \rangle \triangleleft \langle M, < \rangle$.

Für Ordinalzahlen drängt sich nun die folgende Definition einer linearen Ordnung „kleiner als“ auf:

Definition (*die Ordnung auf den Ordinalzahlen*)

Seien α, β Ordinalzahlen. Wir setzen:

$$\alpha < \beta \text{ falls } \langle M, < \rangle \triangleleft \langle N, < \rangle,$$

wobei $\langle M, < \rangle, \langle N, < \rangle$ Wohlordnungen sind mit

$$\alpha = \text{o.t.}(\langle M, < \rangle), \beta = \text{o.t.}(\langle N, < \rangle).$$

Diese Definition hängt nicht von der Wahl von $\langle M, < \rangle$ und $\langle N, < \rangle$ ab.

Offenbar werden die Ordinalzahlen durch $<$ linear geordnet. Wir werden gleich die Wohlordnungseigenschaft von $<$ zeigen: Jede nichtleere Menge A von Ordinalzahlen hat ein kleinstes Element. Tatsächlich hat die Ordnung $<$ auf den Ordinalzahlen die bemerkenswerte Eigenschaft, daß für jede Ordinalzahl α die Menge aller kleineren Ordinalzahlen eine Wohlordnung des Typs α bildet.

Definition (*die Mengen $W(\alpha)$*)

Für Ordinalzahlen α setzen wir:

$$W(\alpha) = \{ \beta \mid \beta \text{ ist Ordinalzahl und } \beta < \alpha \}.$$

Für jedes $n \in \mathbb{N}$ ist $W(n) = \bar{n}$. Weiter gilt $W(\omega) = \mathbb{N} = \omega$.

Satz (*Fundamentalsatz über $W(\alpha)$*)

Sei α eine Ordinalzahl.

Dann ist $\langle W(\alpha), < \rangle$ eine Wohlordnung und es gilt:

$$\text{o.t.}(\langle W(\alpha), < \rangle) = \alpha.$$

Beweis

Sei $\langle M, < \rangle$ eine Wohlordnung des Typs α . Für $x \in M$ sei

$$f(x) = \text{o.t.}(\langle M_x, < \rangle).$$

Dann gilt $f : M \rightarrow W(\alpha)$. Weiter ist f ordnungstreu:

Für alle $x, y \in M$ gilt: $x < y$ gdw $f(x) < f(y)$.

Insbesondere ist f injektiv. f ist aber auch surjektiv:

Sei $\beta < \alpha$, und sei $\langle N, < \rangle$ eine Wohlordnung des Ordnungstyps β .

Dann gilt $\langle N, < \rangle \triangleleft \langle M, < \rangle$ und daher existiert ein $x \in M$ mit $\langle N, < \rangle \equiv \langle M_x, < \rangle$.

Dann gilt aber $f(x) = \text{o.t.}(\langle M_x, < \rangle) = \text{o.t.}(\langle N, < \rangle) = \beta$.

Also ist $f : M \rightarrow W(\alpha)$ bijektiv und ordnungstreu.

Somit ist $\langle W(\alpha), < \rangle$ eine Wohlordnung, und es gilt

$$\text{— } \text{o.t.}(\langle W(\alpha), < \rangle) = \text{o.t.}(\langle M, < \rangle) = \alpha.$$

Hausdorff (1914): „Wir führen hier eine im folgenden durchweg festgehaltene Bezeichnung ein:

$$W(\alpha) = \text{Menge aller Ordnungszahlen } < \alpha.$$

Zunächst zeigen wir, daß jede wohlgeordnete Menge A vom Typus α mit $W(\alpha)$ ähnlich ist ...“

Korollar (*Wohlordnungseigenschaft für die Ordinalzahlen*)

Sei A eine nichtleere Menge von Ordinalzahlen.

Dann besitzt A ein kleinstes Element.

Beweis

Sei $\alpha \in A$ beliebig. Ist α kleinstes Element von A , so sind wir fertig.

Andernfalls ist $B = W(\alpha) \cap A \neq \emptyset$. Da $\langle W(\alpha), < \rangle$ eine Wohlordnung ist, besitzt $B \subseteq W(\alpha)$ ein kleinstes Element und dieses ist dann offenbar

– das kleinste Element von A .

Nachfolger- und Limesordinalzahlen

Für Wohlordnungen hatten wir allgemein Nachfolger- und Limeselemente definiert, sowie die Operationen der Enderweiterung um ein Element und der Vereinigung einer Menge vergleichbarer Wohlordnungen. Diese Begriffe übertragen wir nun auf die Ordinalzahlen.

Definition (*Nachfolger einer Ordinalzahl*)

Sei α eine Ordinalzahl.

Dann ist der *Nachfolger* $\alpha + 1$ von α definiert durch:

$$\alpha + 1 = \text{o. t.}(\langle W(\alpha), < \rangle + \{x\}),$$

wobei x ein Objekt ist mit $x \notin W(\alpha)$.

Wir können hier z. B. $x = \alpha$ wählen.

Definition (*Nachfolgerordinalzahlen und Limesordinalzahlen*)

Sei α eine Ordinalzahl.

- (i) α heißt *Nachfolgerordinalzahl*, falls $\alpha = \beta + 1$ für eine Ordinalzahl β .
In diesem Fall heißt β der *Vorgänger* von α , in Zeichen $\beta = \alpha - 1$.
- (ii) α heißt *Limesordinalzahl*, falls $\alpha \neq 0$ und α keine Nachfolgerordinalzahl ist.

Offenbar ist α eine Nachfolgerordinalzahl, falls α ein Nachfolgerelement in $\langle W(\beta), < \rangle$ ist für ein (alle) $\beta > \alpha$. Analoges gilt für Limesordinalzahlen.

Übung

- (i) $\alpha + 1$ ist die kleinste Ordinalzahl, die größer ist als α .
- (ii) α ist Limesordinalzahl gdw $\alpha \neq 0$ und $\langle W(\alpha), < \rangle$ hat kein größtes Element.

Ist A eine Menge von Ordinalzahlen, so ist $\Gamma = \{ \langle W(\alpha), < \rangle \mid \alpha \in A \}$ eine Menge von Wohlordnungen mit der Eigenschaft: Sind $\langle M, < \rangle$ und $\langle N, < \rangle$ verschiedene Elemente von Γ , so ist $\langle M, < \rangle$ ein Anfangsstück von $\langle N, < \rangle$ oder umgekehrt. Also ist $\bigcup \Gamma$ definiert.

Definition (*Supremum einer Menge von Ordinalzahlen*)

Sei A eine Menge von Ordinalzahlen. Dann ist das *Supremum* von A , in Zeichen $\sup(A)$, definiert durch:

$$\sup(A) = \text{o. t.}(\bigcup \{W(\alpha), <\} \mid \alpha \in A\}).$$

Speziell ist $\sup(\emptyset) = 0$, denn $\text{o. t.}(\langle \emptyset, \emptyset \rangle) = 0$.

Übung

Sei A eine Menge von Ordinalzahlen, und sei $\sigma = \sup(A)$. Dann ist σ die kleinste Ordinalzahl mit $A \leq \sigma$, d. h. es gilt:

- (i) für alle $\alpha \in A$ ist $\alpha \leq \sigma$,
- (ii) ist σ' eine Ordinalzahl mit $\alpha \leq \sigma'$ für alle $\alpha \in A$, so ist $\sigma \leq \sigma'$.

In Kapitel 13 dieses Abschnitts werden wir die bemerkenswerte Eigenschaft, daß sich zu jeder Menge A von Ordinalzahlen eine größere Ordinalzahl finden läßt, etwa $\sup(A) + 1$, genauer betrachten.

Eine zuweilen nützliche Spielerei mit Definitionen ist:

Übung

Sei α eine Ordinalzahl, $\alpha \neq 0$. Dann gilt:

- (i) α ist Limesordinalzahl *gdw* $\sup(W(\alpha)) = \alpha$.
- (ii) Ist $\sup(W(\alpha)) \neq \alpha$, so ist $\sup(W(\alpha)) = \alpha - 1$.

So ist z. B. $\sup(W(17)) = 16$, $\sup(W(\omega)) = \omega$.

Transfinite Folgen und Aufzählungen einer Menge

Die natürlichen Zahlen tauchen häufig als Indizes an bestimmten Objekten auf, etwa bei einer Folge $x_0, x_1, \dots, x_n, \dots$, $n \in \mathbb{N}$, bestehend aus reellen Zahlen. Derartige Folgen sind offiziell nichts anderes als Funktionen $f: \mathbb{N} \rightarrow \mathbb{R}$ mit $f(n) = x_n$ für $n \in \mathbb{N}$. Analog können wir nun Objekte mit Ordinalzahlen indizieren. Dies führt zum Begriff der transfiniten Folge.

Definition (*transfinite Folgen*)

Seien α eine Ordinalzahl, M eine Menge.
 Weiter sei $f: W(\alpha) \rightarrow M$ eine Funktion.
 Dann heißt f eine *Folge der Länge α in M* .
 Ist $\alpha \geq \omega$, so heißt f eine *transfinite Folge* in M .
 Wir schreiben f auch in der Form:

$$f = \langle x_\beta \mid \beta < \alpha \rangle,$$

$$\text{mit } x_\beta = f(\beta) \text{ für } \beta < \alpha.$$

die Schreibweise

$$f = \langle x_\beta \mid \beta < \alpha \rangle$$

für Funktionen auf $W(\alpha)$

Verwandte Schreibweisen, etwa

$\langle x_\beta \mid \alpha' < \beta < \alpha \rangle$ für ein $\alpha' < \alpha$,

$\langle y_\beta \mid \beta \in A \rangle$ für eine Menge A von Ordinalzahlen,

usw. verwenden wir ohne weiteren Kommentar. $\langle y_\beta \mid \beta \in A \rangle$ ist z. B. die Funktion f mit $\text{dom}(f) = A$ und $f(\beta) = y_\beta$ für alle $\beta \in A$.

Transfinite Folgen werden in der Mengenlehre sehr häufig verwendet. Man kann sich eine Folge $\langle x_\beta \mid \beta < \alpha \rangle$ in M vorstellen als eine Reihe (+) wie zu Beginn dieses Abschnitts, wobei nun Elemente aus M die Reihe bilden und die Ordinalzahlen als Indizes verwendet werden:

$$(+) \quad x_0, x_1, x_2, \dots, x_\omega, x_{\omega+1}, \dots, x_{\omega+\omega}, \dots, \dots, x_\beta, \dots, \dots \quad (\beta < \alpha)$$

Ist $f: W(\alpha) \rightarrow M$ sogar bijektiv, so induziert f eine Wohlordnung auf M . Wir können dann f als eine Aufzählung der Elemente von M auffassen. Die entsprechende Reihe (+) durchläuft dann ganz M ohne Wiederholungen, und genauer gilt: Die Elemente von M werden durch diese Reihe in eine Wohlordnung gebracht.

Definition (*Aufzählung und induzierte Wohlordnung*)

Seien α eine Ordinalzahl, M eine Menge und $f: W(\alpha) \rightarrow M$ bijektiv.

Dann heißt $f = \langle x_\beta \mid \beta < \alpha \rangle$ eine *Aufzählung von M* .

Für $x, y \in M$ setzen wir

$$x < y \text{ falls } f^{-1}(x) < f^{-1}(y).$$

$\langle M, < \rangle$ heißt die *durch f induzierte Wohlordnung* auf M .

$\langle M, < \rangle$ ist eine Wohlordnung mit o. t. $(\langle M, < \rangle) = \alpha$ (vgl. auch Kapitel 3).

Ist M schon wohlgeordnet, so können wir eine Aufzählung von M entlang der vorgegebenen Wohlordnung durchführen:

Definition (*Aufzählung einer Wohlordnung*)

Sei $\langle M, < \rangle$ eine Wohlordnung, und sei $\alpha = \text{o. t.}(\langle M, < \rangle)$.

Weiter sei $f: W(\alpha) \rightarrow M$ ordnungsisomorph.

Dann heißt $f = \langle x_\beta \mid \beta < \alpha \rangle$ die *Aufzählung von $\langle M, < \rangle$* .

Für $\beta < \alpha$ heißt x_β das β -te Element von $\langle M, < \rangle$.

Die Beobachtung, daß $\langle W(\alpha), < \rangle$ eine Wohlordnung des Typs α ist, kann man rückblickend zu einer Definition des Begriffs der Ordinalzahl heranziehen. Wir besprechen nun diesen Weg, der zu einer strengen Definition des Ordnungstyps von Wohlordnungen führt, und einen Abstraktionsakt oder ein „denken wir uns zugeordnet ...“ nicht mehr benötigt. Auf diese Weise werden Ordinalzahlen heute üblicherweise eingeführt, wobei die Gefahr besteht, daß die Idee der Ordinalzahl durch die gnadenlose Eleganz der Definition verschüttet wird.

Die moderne Definition einer Ordinalzahl

Das Ziel ist, in jeder Klasse gleichlanger Wohlordnungen einen ausgezeichneten Repräsentanten zu finden. Dies gelingt nun überraschend leicht, wenn man sich von der Idee leiten läßt, die Menge $W(\alpha) = \{\beta \mid \beta < \alpha\}$ mit der Ordinalzahl α selbst zu identifizieren, daß also eine Ordinalzahl nichts anderes ist als die Menge ihrer Vorgänger. Die sich aus dem Postulat „ $W(\alpha) = \alpha$!“ ergebende Definition der Ordinalzahlen stammt von John von Neumann (1923) und davon unabhängig von Ernst Zermelo.

Für diese Definition der Ordinalzahlen entwickeln wir die Theorie der Wohlordnungen wie in Kapitel 3 und 4. (Der Wohlordnungssatz wird, wie für die Cantorsche Definition, nicht benötigt.) Nun setzen wir:

Definition (*Ordinalzahlen nach von Neumann und Zermelo*)

Eine Menge α heißt eine (*Neumann-Zermelo-*) *Ordinalzahl*, falls eine Wohlordnung $<$ auf α existiert mit:

($\diamond$) Für alle $\beta \in \alpha$ ist $\beta = \{\gamma \in \alpha \mid \gamma < \beta\}$.

$\langle \alpha, < \rangle$ heißt dann eine *ordinale Wohlordnung von α* .

Zur Verwendung von griechischen Buchstaben: Wir zeigen gleich, daß alle Elemente einer Ordinalzahl wieder Ordinalzahlen sind.

Schreiben wir wieder kurz α_β für das durch $\beta \in \alpha$ bestimmte Anfangsstück von $\langle \alpha, < \rangle$, d. h. $\alpha_\beta = \{\gamma \in \alpha \mid \gamma < \beta\}$, so lautet ($\diamond$) einfach:

($\diamond$) $\beta = \alpha_\beta$ für alle $\beta \in \alpha$.

Unmittelbar aus der Definition fließen einige wichtige Eigenschaften der Ordinalzahlen:

Satz (*elementare Eigenschaften von Ordinalzahlen*)

Sei α eine Ordinalzahl, und sei $\langle \alpha, < \rangle$ eine ordinale Wohlordnung von α . Dann gilt:

- (i) Für alle $\beta, \gamma \in \alpha$ gilt: $\gamma < \beta$ gdw $\gamma \in \beta$ gdw $\gamma \subset \beta$.
- (ii) Ist $\beta \in \alpha$, so ist $\beta \subseteq \alpha$.
- (iii) Ist $\beta \in \alpha$, so ist β eine Ordinalzahl.

Beweis

zu (i): Seien $\beta, \gamma \in \alpha$. Es gilt: $\gamma < \beta$ gdw $\gamma \in \alpha_\beta$ gdw $\gamma \in \beta$.

Weiter gilt: $\gamma < \beta$ gdw $\alpha_\gamma \subset \alpha_\beta$ gdw $\gamma \subset \beta$.

zu (ii): Sei $\beta \in \alpha$. Dann gilt $\beta = \alpha_\beta \subseteq \alpha$.

zu (iii): Sei $\beta \in \alpha$, und sei $<_\beta = \{(\alpha_1, \alpha_2) \mid \alpha_1, \alpha_2 \in \alpha, \alpha_1 < \alpha_2 < \beta\}$.

Wir zeigen, daß $\langle \beta, <_\beta \rangle$ eine ordinale Wohlordnung ist.

$\langle \beta, <_\beta \rangle$ ist eine Wohlordnung wegen $\beta \subseteq \alpha$.

Sei nun $\gamma \in \beta$. Dann gilt:

$$\beta_\gamma = \{ \delta \in \beta \mid \delta <_\beta \gamma \} = \beta \cap \alpha_\gamma = \beta \cap \gamma = \gamma,$$

– letzteres wegen $\gamma \subseteq \beta$ nach (i).

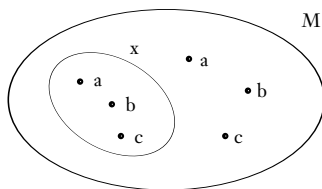
Ordinalzahlen bestehen also ausschließlich aus Ordinalzahlen!

Weiter gilt nach (i): Eine zu einer Ordinalzahl α gehörige ordinale Wohlordnung ist eindeutig bestimmt: Sind $\langle \alpha, <_1 \rangle$ und $\langle \alpha, <_2 \rangle$ ordinale Wohlordnungen, so gilt $<_1 = <_2$. Ist α eine Ordinalzahl, so werden die Ordinalzahlen, die kleiner als α sind, sowohl durch die $\in$ -Relation auf α als auch durch die $\subset$ -Relation auf α wohlgeordnet – und beide Relationen ergeben die gleiche Wohlordnung.

Eigenschaft (ii) spielt in der Mengenlehre an verschiedenen Stellen eine wichtige Rolle, und hat deswegen einen eigenen Namen:

Definition (*transitiv*)

Sei M eine Menge. M heißt *transitiv*, falls für alle $x \in M$ gilt $x \subseteq M$.



Transitivität sollte jedem vernünftigen Modell M der Mengenlehre zukommen: Ist x ein Element des Modells, so bedeutet $x \subseteq M$, daß jedes Element von x ein Objekt des Modells ist.

Übung

Sei M eine Menge. Dann gilt:

M ist transitiv gdw $\bigcup M \subseteq M$ gdw $M \subseteq \mathcal{P}(M)$.

Ordinalzahlen sind nun charakterisiert durch die Bedingungen „transitiv“ und „durch $\in$ wohlgeordnet“, was manchmal auch zur Definition von Ordinalzahl benutzt wird:

Satz (*Charakterisierung der Ordinalzahlen*)

Sei α eine Menge. Dann sind äquivalent:

- (i) α ist eine Ordinalzahl.
- (ii) α ist transitiv und wird durch $\in$ wohlgeordnet,
d.h. $\langle \alpha, \in \upharpoonright \alpha \rangle = \langle \alpha, \{ (\beta, \gamma) \in \alpha \times \alpha \mid \beta \in \gamma \} \rangle$ ist eine Wohlordnung.

Beweis

(i) $\curvearrowright$ (ii): Folgt aus dem Satz oben.

(ii) $\curvearrowright$ (i): Sei $\beta \in \alpha$. Wegen α transitiv ist dann $\beta \subseteq \alpha$, und es gilt

– $\alpha_\beta = \{ \gamma \in \alpha \mid \gamma < \beta \} = \{ \gamma \in \alpha \mid \gamma \in \beta \} = \alpha \cap \beta = \beta$.

Als nächstes zeigen wir, daß gleichlange ordinale Wohlordnungen zusammenfallen:

Satz (*Eindeutigkeitssatz*)

Seien α, β Ordinalzahlen. Weiter seien die zugehörigen ordinalen Wohlordnungen $\langle \alpha, < \rangle$ und $\langle \beta, < \rangle$ gleichlang.
Dann gilt $\alpha = \beta$.

Beweis

Sei $f: \alpha \rightarrow \beta$ ordnungsisomorph. *Annahme* $f \neq \text{id}_\alpha$. Dann existiert ein kleinstes $\gamma \in \alpha$ mit $f(\gamma) \neq \gamma$. Nach minimaler Wahl von γ ist dann $\alpha_\gamma = \beta_{f(\gamma)}$,
– und folglich $\gamma = f(\gamma)$, da α, β Ordinalzahlen sind. *Widerspruch!*

Wir definieren nun die Ordnung auf den Ordinalzahlen:

Definition (*die $\in$ -Ordnung auf den Ordinalzahlen*)

Seien α, β Ordinalzahlen. Wir setzen:

$\alpha < \beta$ falls $\alpha \in \beta$.

Für alle Ordinalzahlen α gilt also

$$\alpha = \{ \beta \mid \beta \in \alpha \} = \{ \beta \mid \beta \in \alpha \text{ und } \beta \text{ ist Ordinalzahl} \} = \{ \beta \mid \beta < \alpha \}.$$

Satz

Die Ordinalzahlen werden durch $<$ wohlgeordnet.

Beweis

$<$ ist *irreflexiv*:

Annahme $\alpha < \alpha$ für eine Ordinalzahl α . Dann gilt $\alpha \in \alpha$, also existiert ein $\beta \in \alpha$ mit $\beta \in \beta$ (α ist ein solches β). Aber $\in$ ist eine Wohlordnung auf α , und damit insbesondere irreflexiv. *Widerspruch!*

$<$ ist *transitiv*:

Seien α, β, γ Ordinalzahlen mit $\alpha < \beta$ und $\beta < \gamma$.

Dann gilt $\alpha \in \beta \in \gamma$. Wegen γ transitiv ist dann $\alpha \in \gamma$, also $\alpha < \gamma$.

$<$ ist *linear*:

Seien α, β Ordinalzahlen. Sind $\langle \alpha, < \rangle$ und $\langle \beta, < \rangle$ gleichlang, so gilt $\alpha = \beta$ nach dem Satz oben.

Ist $\langle \alpha, < \rangle$ kürzer als $\langle \beta, < \rangle$, so existiert ein $\gamma \in \beta$ derart, daß $\langle \alpha, < \rangle$ und $\langle \gamma, < \rangle$ gleichlang sind, und dann gilt $\alpha = \gamma$, also $\alpha \in \beta$. Analog folgt $\beta \in \alpha$, falls $\langle \beta, < \rangle$ kürzer als $\langle \alpha, < \rangle$ ist.

Wohlordnungseigenschaft:

Sei A eine nichtleere Menge von Ordinalzahlen, und sei $\alpha \in A$.

Ist α kleinstes Element von A , so sind wir fertig.

Andernfalls ist $\alpha \cap A$ nichtleer.

Sei also β das kleinste Element von $\alpha \cap A$ in der Wohlordnung $\langle \alpha, < \rangle$.

Dann ist β das kleinste Element von A :

Denn für alle $\gamma \in \alpha \cap A$ ist $\beta \leq \gamma$ nach Wahl von β , und für

– $\gamma \in A - \alpha$ ist $\alpha \leq \gamma$, also $\beta < \alpha \leq \gamma$.

Elementare Fakten über $<$ sind:

Übung

- (i) Für jede Ordinalzahl α ist $\alpha \cup \{\alpha\}$ die kleinste Ordinalzahl, die größer ist als α .
- (ii) Ist A eine Menge von Ordinalzahlen, so ist $\sigma = \bigcup A$ eine Ordinalzahl. Weiter ist $\bigcup A$ das kleinste α mit $A \leq \alpha$.
- (iii) Ist A eine nichtleere Menge von Ordinalzahlen, so ist $\tau = \bigcap A$ das Minimum von A , d.h. $\tau \in A$ und $\tau \leq A$.

Hieraus ergeben sich die folgenden Definitionen:

Definition (Nachfolger einer Ordinalzahl)

Sei α eine Ordinalzahl. Dann ist der *Nachfolger* $\alpha + 1$ von α definiert durch:

$$\alpha + 1 = \alpha \cup \{\alpha\}.$$

Definition (Nachfolgerordinalzahlen und Limesordinalzahlen)

Sei α eine Ordinalzahl.

- (i) α heißt *Nachfolgerordinalzahl*, falls $\alpha = \beta + 1$ für eine Ordinalzahl β . In diesem Fall heißt β der *Vorgänger* von α , in Zeichen $\beta = \alpha - 1$.
- (ii) α heißt *Limesordinalzahl*, falls $\alpha \neq 0$ und α keine Nachfolgerordinalzahl ist.

In vielen Texten wird auch $\emptyset$ als Limesordinalzahl angesehen, sodaß jede Ordinalzahl Nachfolger oder Limes ist.

Definition (Supremum einer Menge von Ordinalzahlen)

Sei A eine Menge von Ordinalzahlen. Dann ist das *Supremum* von A , in Zeichen $\sup(A)$, definiert durch:

$$\sup(A) = \bigcup A.$$

Wieder gilt $\sup(\emptyset) = \emptyset$.

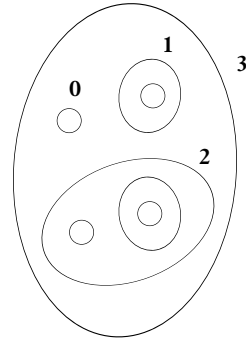
Übung

Sei $\alpha \neq \emptyset$ eine Ordinalzahl. Dann sind äquivalent:

- (i) α ist eine Limesordinalzahl,
- (ii) $\bigcup \alpha = \alpha$,
- (iii) $\beta + 1 \in \alpha$ für alle $\beta < \alpha$.

Wir können uns damit einen Überblick über die ersten Ordinalzahlen verschaffen: Der Nachfolger $\alpha + 1$ von α ist $\alpha \cup \{\alpha\}$, und das Supremum einer Menge von Ordinalzahlen wird einfach durch die Vereinigung über die Menge gegeben. Bezeichnen wir die ersten Ordinalzahlen wieder mit natürlichen Zahlen, so ergibt sich:

$$\begin{aligned}
0 &= \emptyset, \\
1 &= 0 \cup \{0\} = \{\emptyset\} = \{0\}, \\
2 &= 1 \cup \{1\} = \{\emptyset, \{\emptyset\}\} = \{0, 1\}, \\
3 &= 2 \cup \{2\} = \{\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\}\} = \{0, 1, 2\}, \\
&\dots \\
n+1 &= n \cup \{n\} = \{0, 1, \dots, n\}, \\
&\dots \\
\omega &= \mathbb{N} = \{0, 1, 2, \dots\}, \\
\omega+1 &= \omega \cup \{\omega\} = \{0, 1, 2, \dots, \omega\}. \\
&\dots \\
\omega+\omega &= \{0, 1, 2, \dots, \omega, \omega+1, \omega+2, \dots\}. \\
&\dots
\end{aligned}$$



Ostern à la von Neumann

Lösen wir die Elemente einer Ordinalzahl immer weiter auf, so sehen wir, daß alle Ordinalzahlen letztendlich aus der leeren Menge aufgebaut sind. Es zeigt sich darüber hinaus, daß man zur Definition der Ordinalzahlen keinerlei Zahlen als Grundobjekte benötigt. Man kann also, wenn man möchte, die natürlichen Zahlen definieren als die Menge aller Ordinalzahlen, die kleiner sind als die erste Limesordinalzahl, ω genannt, und dadurch wird dann $\mathbb{N} = \omega$. In der axiomatischen Mengenlehre, die keine Grundobjekte zuläßt, ist diese Vorgehensweise üblich. Man zeigt aus den Axiomen die Existenz einer Limesordinalzahl, und definiert dann die kleinste Limesordinalzahl als die Menge der natürlichen Zahlen.

Es bleibt nun noch die Frage zu klären, ob der Begriff der Ordinalzahlen reichhaltig genug ist, um für jede Wohlordnung einen Repräsentanten gleicher Länge zur Verfügung zu stellen. Man sieht leicht, daß die Ordinalzahlen dies auf jeden Fall für einen Anfangsbereich der Wohlordnungen leisten, d. h. bis zu einer bestimmten Länge gibt es auf jeden Fall ordinale Repräsentanten: Denn sind $\langle M, < \rangle$ und $\langle \alpha, < \rangle$ gleichlang, so existiert für alle Wohlordnungen $\langle N, < \rangle$, die kürzer als $\langle M, < \rangle$ sind, ein Anfangsstück $\langle \alpha_\beta, < \rangle$ von $\langle \alpha, < \rangle$, das gleichlang mit $\langle N, < \rangle$ ist. Aber wegen $\alpha_\beta = \beta$ existiert dann also sogar eine ordinale Wohlordnung der Länge von $\langle N, < \rangle$. Die Annahme, daß eine Wohlordnung $\langle M, < \rangle$ in ihrer Länge über alle Ordinalzahlen hinausragt, führt aber zu einem Widerspruch. Den Beweis gibt der folgende Repräsentationssatz für Neumann-Zermelo-Ordinalzahlen. (Für die Definition nach Cantor und Hausdorff gilt der Satz automatisch.)

Satz (*Repräsentationssatz*)

Zu jeder Wohlordnung $\langle M, < \rangle$ existiert eine eindeutige Ordinalzahl α mit:
 $\langle M, < \rangle$ und $\langle \alpha, < \rangle$ sind gleichlang.

Beweis

zur Eindeutigkeit: Dies folgt aus dem Eindeutigkeitssatz oben.

zur Existenz: Annahme nicht für eine Wohlordnung $\langle M, < \rangle$. Wir setzen

$$M' = \{x \in M \mid \text{es existiert eine Ordinalzahl } \alpha \text{ mit } \langle M_x, < \rangle \equiv \langle \alpha, < \rangle\}.$$

Für $x \in M'$ sei $f(x)$ die eindeutige Ordinalzahl α mit $\langle M_x, < \rangle \equiv \langle \alpha, < \rangle$.

Dann ist $f : M' \rightarrow \text{rng}(f)$ bijektiv.

Sei $A = \text{rng}(f)$. Dann ist A eine Menge von Ordinalzahlen und es gilt:

(+) Jede Ordinalzahl α ist ein Element von A .

Beweis von (+)

Denn ist $\alpha \notin A = \text{rng}(f)$, so gilt nach dem Vergleichbarkeitssatz und der Annahme, daß $\langle M, < \rangle \triangleleft \langle \alpha, < \rangle$. Dann existiert aber ein $\beta \in \alpha$ mit $\langle M, < \rangle \equiv \langle \alpha_\beta, < \rangle$. Aber $\alpha_\beta = \beta$, also sind $\langle M, < \rangle$ und $\langle \beta, < \rangle$ gleichlang, *im Widerspruch* zur Annahme. Dies zeigt (+).

Also ist A eine Menge von Ordinalzahlen, und zugleich ist jede Ordinalzahl α ein Element von A . Insbesondere gilt $\text{sup}(A) + 1 \in A$,
 — was für alle Mengen von Ordinalzahlen falsch ist. *Widerspruch!*

Daß $A = \text{rng}(f)$ eine Menge, eine konsistente Vielheit, ist, können wir so begründen: Inkonsistente Vielheiten entstehen durch zu große Zusammenfassungen. Hier ist aber die Zusammenfassung beschränkt, da A die Größe der Menge M' hat: unter der gemachten Annahme ist $f : M' \rightarrow A$ bijektiv. Innerhalb der axiomatischen Mengenlehre wird die Existenz der Menge A in solchen Fällen durch ein spezielles Axiom, das sog. Ersetzungsaxiom, garantiert. In der originalen Zermelo-Axiomatik ohne dieses Axiom ist der Repräsentationssatz nicht beweisbar: man kann dort nicht zeigen, daß es ausreichend viele von-Neumann-Zermelo-Ordinalzahlen gibt, um alle Wohlordnungen erreichen zu können.

Die Ordinalzahlen sind also in so großer Anzahl vorhanden, daß sie Repräsentanten für Wohlordnungen beliebiger Länge liefern. Andererseits sind sie so dünn gesät, daß diese Repräsentanten eindeutig sind. Diese Eigenschaften erlauben uns nun eine lupenreine Definition des Ordnungstyps einer Wohlordnung:

Definition (*Ordnungstyp einer Wohlordnung*)

Sei $\langle M, < \rangle$ eine Wohlordnung. Dann ist der *Ordnungstyp* von $\langle M, < \rangle$, in Zeichen o.t. $\langle \langle M, < \rangle \rangle$ definiert durch:

o.t. $\langle \langle M, < \rangle \rangle =$ „die Ordinalzahl α mit:

$\langle M, < \rangle$ und $\langle \alpha, < \rangle$ sind gleichlang“.

Die Definition der Ordinalzahlen in dieser Sektion ist auf den ersten Blick schwer verdaulich und greift der Zeit der Mengenlehre Cantors auch voraus. Es zeigt sich aber schnell, daß man mit ihr überaus bequem arbeiten kann. Die *Idee* einer Ordinalzahl zu entwickeln bleibt jedem selbst überlassen, und ob man sich dabei eher am Cantorschen Abstraktionsprozeß oder an der modernen und formal durchführbaren Definition orientiert, ist Geschmackssache.

von Neumann (1923): „Das Ziel der vorliegenden Arbeit ist: den Begriff der *Cantorschen* Ordnungszahl eindeutig und konkret zu fassen.

Dieser Begriff wird nach *Cantors* Vorgang gewöhnlich als ‚Abstraktion‘ einer gemeinsamen Eigenschaft aus gewissen Klassen gewonnen. Dieses etwas vage Verfahren wollen wir durch ein anderes auf eindeutigen Mengenoperationen beruhendes, ersetzen. Das Verfahren wird in den folgenden Zeilen in der Sprache der naiven Mengen-

lehre dargestellt werden, es bleibt aber (im Gegensatz zu Cantors Verfahren) auch in einer ‚formalistischen‘, axiomatisierten Mengenlehre richtig. So behalten unsere Schlüsse auch im Rahmen der Zermeloschen Axiomatik (wenn man das Fraenkelsche Axiom [Ersetzungsaxiom] hinzufügt) volle Geltung.

Wir wollen eigentlich den Satz: ‚Jede Ordnungszahl ist der Typus der Menge aller ihr vorangehenden Ordnungszahlen‘ zur Grundlage unserer Überlegungen machen. Damit aber der vage Begriff ‚Typus‘ vermieden werde, in dieser Form: ‚Jede Ordnungszahl ist die Menge der ihr vorangehenden Ordnungszahlen.‘“

Der Leser kann im folgenden die Definition von Cantor-Hausdorff zugrundelegen oder die Neumann-Zermelo-Ordinalzahlen verwenden. Wir schreiben weiterhin $W(\alpha) = \{\beta \mid \beta < \alpha\}$. Für Neumann-Zermelo-Ordinalzahlen ist dann zudem an jeder Stelle einfach $W(\alpha) = \alpha$ (und $\alpha < \beta$ ist nichts anderes als $\alpha \in \beta$, und darüber hinaus identisch mit $\alpha \subset \beta$).

Mächtigkeiten und Kardinalzahlen

Wir definieren nun mit Hilfe der Ordinalzahlen die Kardinalzahlen, und anschließend, mit Hilfe des Wohlordnungssatzes, die Mächtigkeit $|M|$ einer Menge. Wir hatten im ersten Abschnitt Cantors Definition von „Mächtigkeit“ und „Kardinalzahl“ von 1887 angegeben (Ende von 1.4), aber erst in 1.12 mit Kardinalzahlen nach Cantor und Hausdorff gerechnet. Wir können nun die definitorische Hypothek des Kapitels über Kardinalzahlarithmetik mit Hilfe der Hypothek der Cantor-Hausdorff-Ordinalzahlen zurückzahlen, oder unter Verwendung der Neumann-Zermelo-Ordinalzahlen völlig schuldenfrei werden.

In Cantors Arbeit von 1895 findet man gleich nach seiner Definition von „Menge“ die folgende doppelte Abstraktion zu Einsen oder Einheiten:

Cantor (1895): „Mächtigkeit‘ oder ‚Kardinalzahl‘ von M nennen wir den Allgemeinbegriff, welcher mit Hilfe unseres aktiven Denkvermögens dadurch aus der Menge M hervorgeht, daß von der Beschaffenheit ihrer verschiedenen Elemente m und von der Ordnung ihres Gegebenseins abstrahiert wird.

Das Resultat dieses zweifachen Abstraktionsakts, die Kardinalzahl oder Mächtigkeit von M , bezeichnen wir mit

$$(3) \quad \overline{\overline{M}}.$$

Da aus jedem einzelnen Elemente m , wenn man von seiner Beschaffenheit absieht, eine ‚Eins‘ wird, so ist die Kardinalzahl $\overline{\overline{M}}$ selbst eine bestimmte aus lauter Einsen zusammengesetzte Menge, die als intellektuelles Abbild oder Projektion der gegebenen Menge M in unserem Geiste Existenz hat.“

Cantors Sicht ist also diese: Sieht man von der Natur der Elemente von M ab, so bleibt zunächst noch eine evtl. vorhandene Ordnung von M übrig. Löst man auch diese Ordnung auf, so bleibt ein Gebilde aus ununterscheidbaren

Objekten als Ergebnis. Jede andere Menge N , für die eine Bijektion von N nach M existiert, liefert genau das gleiche Gebilde, die zu beiden Mengen gehörige Kardinalzahl. Cantor hat Funktionen und Ordnungen nicht als Mengen betrachtet, und ebenso liegen die durch Abstraktion gewonnenen Ordinal- und Kardinalzahlen außerhalb der „heutigen“, extensionalen Mengenwelt. (Andernfalls wäre nach dem Extensionalitätsaxiom die Kardinalität einer nichtleeren Menge unter Cantors Definition immer identisch mit $\{1\}$.) Suggestiv könnte man eine Cantorsche Ordinalzahl als $\langle\langle 1, \dots, 1, \dots, 1, \dots \rangle\rangle$ schreiben, wobei die Folge der Einsen eine Wohlordnung einer bestimmten Länge darstellt. Und weiter wäre dann eine Kardinalzahl eine Multimenge $\{\{1, \dots, 1, \dots, 1, \dots\}\}$, in der es zwar nicht mehr auf die Ordnung der Einsen untereinander, aber auf ihre Vielfachheit ankommt.

Die Cantorsche Definition durch Abstraktion hat eine lange Tradition. In der griechischen Mathematik gibt es den Begriff der Einheit oder Monade. Natürliche Zahlen werden mit dieser Hypothek als Systeme von Einheiten aufgefaßt. (Zur Zahl als $\mu\upsilon\nu\acute{\alpha}\delta\omicron\nu\sigma\acute{\upsilon}\sigma\eta\mu\alpha$ bei Thales, Nikomachos von Gerasa und der „Zahl als zusammengesetzte Vielheit, bestehend aus Einheiten“ bei Euklid vgl. 1.1.) Die griechischen Ideen werden dann dem Mittelalter etwa durch Cassiodorus (~ 477 – 565) und Boethius (~ 480 – 524) zur weiteren Bearbeitung übergeben. Boethius notiert „*numerus est unitatum collectio*“, Cassiodorus „*numerus est ex monadibus multitudo composita*“. Im Mittelalter sind dann derlei „*numerus est*“-Definitionen Folklore. Wie sehr auch noch Cantor von diesen Gedanken überzeugt war, zeigt der Zusatz zu seiner Definition von „Mächtigkeit oder Kardinalzahl“ durch Abstraktion: Elemente werden zu Einsen. Cantors Einsen dürfen also als eine Version der guten alten griechischen Monaden gelten. (Zur Geschichte der Zahl siehe etwa [Gericke 1970, 1971, 1973].)

Heute definiert man allgemein Kardinalzahlen als bestimmte Ordinalzahlen, und mit Hilfe der Neumann-Zermelo-Definition einer Ordinalzahl hat man dann diese zentralen Konstrukte der Mengenlehre innerhalb der extensional iterativen Mengenwelt selbst interpretiert, was ein weiteres Beispiel dafür darstellt, daß sich jede mathematische Theorie, auch die Cantorsche Mengenlehre, unter dem Dach der modernen (formal-axiomatischen) Mengenlehre entwickeln läßt.

Wir definieren also:

Definition (*Kardinalzahl*)

Sei α eine Ordinalzahl. α heißt *Kardinalzahl*, falls gilt:

Für alle $\beta < \alpha$ ist $|W(\beta)| < |W(\alpha)|$.

Die mit Vorliebe verwendeten Zeichen für Kardinalzahlen sind κ , λ und μ .

Die Kardinalzahlen sind die Sprungstellen in der Reihe der Ordinalzahlen hinsichtlich der Existenz von Bijektionen. α ist Kardinalzahl, falls sich für alle $\beta < \alpha$ die Menge $W(\beta)$ nicht bijektiv auf $W(\alpha)$ abbilden läßt.

Obige Definition benutzt nur die Relation $|M| < |N|$, und setzt keine Definition der Mächtigkeit $|M|$ von M selbst voraus. Wir können nun aber mit Hilfe des Wohlordnungssatzes $|M|$ in unaffektierter Weise als Kardinalzahl definieren:

Definition (*Mächtigkeit einer Menge*)

Sei M eine Menge. Dann ist die *Mächtigkeit* oder *Kardinalität* von M , in Zeichen $|M|$, definiert durch:

$|M| =$ „die kleinste Ordinalzahl α mit: $|M| = |W(\alpha)|$ “.

Eine Ordinalzahl β mit $|M| = |W(\beta)|$ existiert nach dem Wohlordnungssatz: Ist $\langle M, < \rangle$ eine Wohlordnung von M und $\beta = o. t. \langle \langle M, < \rangle \rangle$, so gilt $|M| = |W(\beta)|$. Dann existiert aber auch eine kleinste Ordinalzahl α mit $|M| = |W(\alpha)|$, denn $A = \{ \alpha \mid \alpha \leq \beta \text{ und } |M| = |W(\alpha)| \}$ ist eine nichtleere Menge von Ordinalzahlen und besitzt daher ein kleinstes Element.

Die Ordinalzahlen enthalten also mit den Kardinalzahlen eine Meßlatte für die Größe einer Menge, die den Zahlcharakter des Mächtigkeitsbegriffs besonders deutlich macht.

Übung

- (i) Für alle Mengen M ist $|M|$ eine Kardinalzahl.
- (ii) Eine Ordinalzahl α ist genau dann eine Kardinalzahl, wenn $|W(\alpha)| = \alpha$ gilt.

Schließlich ordnen wir noch den Ordinalzahlen selbst eine Kardinalität zu.

Definition (*Kardinalität einer Ordinalzahl*)

Für eine Ordinalzahl α schreiben wir kurz $|\alpha|$ für $|W(\alpha)|$.

$|\alpha|$ heißt die *Kardinalität* von α .

Diese Definition ist unter der modernen Definition einer Ordinalzahl leer, da dann $W(\alpha) = \alpha$ gilt.

Kardinalzahlarithmetik, Nachfolger und Suprema

Die Arithmetik mit Kardinalzahlen können wir nun wie in Kapitel 1.12 entwickeln, wobei sich vieles vereinfacht, da nun die Kardinalzahlen $\alpha = \kappa$ mit einer Wohlordnung versehen sind. Die arithmetischen Operationen definieren wir etwas direkter als früher wie folgt:

Definition (*Addition, Multiplikation und Exponentiation von Kardinalzahlen*)

Seien κ und λ Kardinalzahlen. Wir definieren die *Summe* $\kappa + \lambda$, das *Produkt* $\kappa \cdot \lambda$ und die *Exponentiation* κ^λ von κ und λ wie folgt:

$$\kappa + \lambda = |W(\kappa) \times \{0\} \cup W(\lambda) \times \{1\}|,$$

$$\kappa \cdot \lambda = |W(\kappa) \times W(\lambda)|,$$

$$\kappa^\lambda = |^{W(\lambda)}W(\kappa)|.$$

Der Leser möge sich nun die Kardinalzahlarithmetik aus 1.12 in Erinnerung rufen mit Ausnahme aller Argumente, die Zermelosysteme benutzen. Die ande-

ren Beweise gelten wörtlich, wobei wir nun griechische Buchstaben statt α , β , γ verwenden, denn Kardinalzahlen sind nun spezielle Ordinalzahlen. Die Argumente mit Zermelosystemen werden nun trivial mit Ausnahme des Multiplikationssatzes, den wir später noch einmal beweisen. Insbesondere bekommen wir nun durch die Einbettung der Kardinalzahlen in das ordinale Rückgrat der Mengenwelt die Existenz von Nachfolgerkardinalzahlen und die Existenz von Suprema von Mengen von Kardinalzahlen geschenkt.

Der Leser, der 1.12 nicht ausgelassen hat, wird rückblickend erkennen, daß wir dort die wohlgeordnete Struktur der Kardinalzahlen bereits gezeigt haben, und daß nun die Kardinalzahlen ein Abbild der Ordinalzahlen innerhalb der Ordinalzahlen selbst bilden. Im nächsten Kapitel kommen wir darauf noch zurück.

Zu jeder Kardinalzahl κ finden wir eine größere Kardinalzahl, z. B. 2^κ . Also können wir definieren:

Definition (*kardinaler Nachfolger*)

Sei κ eine Kardinalzahl. Dann ist der (*kardinale*) *Nachfolger* von κ , in Zeichen κ^+ , definiert durch:

$\kappa^+ =$ „die kleinste Kardinalzahl μ mit $\kappa < \mu$ “.

Die Kontinuumshypothese und ihre Verallgemeinerung können wir nun prägnant schreiben als:

(CH) Es gilt $2^\omega = \omega^+$.

(GCH) Für alle unendlichen Kardinalzahlen κ gilt $2^\kappa = \kappa^+$.

Die auf ω folgende Kardinalzahl ist in der Mengenlehre von zentraler Bedeutung und verdient eine eigene Definition:

Definition (ω_1)

Wir setzen $\omega_1 = \omega^+$.

Eine Ordinalzahl α heißt abzählbar, falls $W(\alpha)$ abzählbar ist. Dann gilt:

$\omega_1 =$ „die kleinste überabzählbare Ordinalzahl“,
 $=$ „die kleinste überabzählbare Kardinalzahl“.

Neben ω_1 ist auch die Cantorsche Aleph-Bezeichnung $\aleph_1$ [gelesen: Aleph 1] gebräuchlich. Weiter haben wir $\aleph_0 = \omega_0 = \omega = \aleph$. Es gilt $\aleph_1 = \aleph_0^+$.

Für Cantor bezeichnen Omegas und Alephs verschiedene Dinge: $\aleph_0 = \overline{\aleph}$ ist die Kardinalität der natürlichen Zahlen im Sinne einer „zweifachen Abstraktion“ von der Natur der Elemente und ihrer Ordnung; dagegen ist $\omega = \aleph$ der Ordnungstypus der natürlichen Zahlen im Sinne einer einfachen Abstraktion von der Natur der Elemente unter Beibehaltung ihrer Ordnung. Cantor sieht $\aleph_0$ auch als die Kardinalzahl des Ordnungstypus ω an, was aus seiner Definition der Kardinalzahl einer Menge nicht hervorgeht, was er aber später explizit festhält (1895, § 7). Allgemein ist für eine Ordinalzahl α dann $\overline{\alpha}$ die Kardinalität von α , wobei der Einzelstrich die lediglich noch fehlende einfache Abstraktion von der Ordnung anzeigt. Für Cantor ist dann $\aleph_1 = \overline{\omega_1}$.

In der Mengenlehre werden zuweilen axiomatische Theorien untersucht, in denen nur gewisse Mengen wohlordenbar sind. In diesem Fall verwendet man Alephs für die Kardinalitäten von wohlordenbaren Mengen, und Frakturbuchstaben α, β, γ für beliebige Kardinalzahlen.

Die Kontinuumshypothese lautet nun:

(CH) Es gilt $2^{\omega} = \omega_1$.

$$2^{\omega} = \omega_1$$

Übung

Es gilt $\omega_1^{\omega} = 2^{\omega}$.

Cantors Kontinuumshypothese

$[\omega_1^{\omega} \leq (2^{\omega})^{\omega} = 2^{\omega \cdot \omega} = 2^{\omega}.]$

Wir können die Hypothese auch schreiben als $|\mathbb{R}| = \omega_1$.

Eine vielleicht nützliche Vorstellung ist: ω_1 ist winzig, wenn wir es als ein Mitglied aller Kardinalzahlen betrachten. ω_1 ist dagegen weit weg, wenn wir versuchen, es durch Zählen $0, 1, 2, \dots, \omega, \omega + 1, \dots$ zu erreichen. Das Supremum einer Reihe von abzählbaren Ordinalzahlen $\alpha_0 < \alpha_1 < \dots < \alpha_n < \dots, n \in \mathbb{N}$, ist eine abzählbare Ordinalzahl, kann also nie ω_1 sein. Wir brauchen also überabzählbar viele Schritte, um ω_1 von unten zu erreichen.

Weiter setzen wir $\aleph_2 = \omega_2 = \omega_1^+$, und allgemein $\aleph_{n+1} = \omega_{n+1} = \omega_n^+$ für natürliche Zahlen n . Dann bilden also

$\omega_0, \omega_1, \omega_2, \dots, \omega_n, \dots$

ein Anfangsstück der Reihe der unendlichen Kardinalzahlen. Der Limes dieser Folge ist wieder eine Kardinalzahl. Wir zeigen hierzu allgemein, daß das ordinale Supremum einer Menge von Kardinalzahlen wieder eine Kardinalzahl ist:

Satz (*Suprema von Kardinalzahlen*)

Sei A eine nichtleere Menge von Kardinalzahlen.

Dann ist $\sup(A)$ eine Kardinalzahl.

Beweis

Die Aussage ist klar, falls $\sup(A) \in A$. Sei also $\mu = \sup(A) \notin A$.

Annahme, μ ist keine Kardinalzahl. Dann existiert ein $\alpha < \mu$ mit

$|W(\alpha)| = |W(\mu)|$. Wegen $\mu = \sup(A)$ existiert nun ein $\kappa \in A$ mit $\alpha < \kappa$.

Trivialerweise gilt $|W(\alpha)| \leq |W(\kappa)| \leq |W(\mu)|$. Mit $|W(\alpha)| = |W(\mu)|$

– folgt $|W(\alpha)| = |W(\kappa)|$, im Widerspruch zu $\alpha < \kappa$ und κ Kardinalzahl.

Die auf die Reihe $\omega_1, \omega_2, \dots, \omega_n, \dots$ als nächstes folgende Kardinalzahl wird mit $\aleph_{\omega}$ oder ω_{ω} bezeichnet:

$\aleph_{\omega} = \omega_{\omega} = \sup \{ \omega_n \mid n \in \mathbb{N} \}.$

Wir definieren schließlich noch:

Definition (*Nachfolgerkardinalzahl, Limeskardinalzahl*)

Eine Kardinalzahl $\mu > 0$ heißt *Nachfolgerkardinalzahl*, falls eine Kardinalzahl κ existiert mit $\mu = \kappa^+$. Andernfalls heißt μ eine *Limeskardinalzahl*.

Nach dem Satz ist eine Limeskardinalzahl μ das Supremum aller Kardinalzahlen $\kappa < \mu$.

Cantor (1895): „Aus $\aleph_0$ geht nach einem bestimmten Gesetze die *nächstgrößere* Kardinalzahl $\aleph_1$, aus dieser nach demselben Gesetze die *nächstgrößere* $\aleph_2$ hervor und so geht es weiter.“

Aber auch die unbegrenzte Folge der Kardinalzahlen

$$\aleph_0, \aleph_1, \aleph_2, \dots, \aleph_\nu, \dots$$

erschöpft nicht den Begriff der transfiniten Kardinalzahl. Es wird die Existenz einer Kardinalzahl nachgewiesen werden, die wir mit $\aleph_\omega$ bezeichnen und welche sich als die zu *allen* $\aleph_\nu$ *nächstgrößere* ausweist; aus ihr geht in derselben Weise wie $\aleph_1$ aus $\aleph_0$ eine nächstgrößere $\aleph_{\omega+1}$ hervor und so geht es ohne Ende fort.“

$\aleph_\omega$ hat die bemerkenswerte Eigenschaft, das Supremum einer abzählbaren Menge von Kardinalzahlen zu sein. Dennoch ist $\aleph_\omega$ eine Kardinalzahl und größer als jedes $\aleph_n$ für $n \in \mathbb{N}$.

Die gesamte Aleph-Reihe der unendlichen Kardinalzahlen definieren wir im nächsten Kapitel, wo wir allgemein Induktionen und Rekursionen entlang der Ordinalzahlen – oder entlang von Wohlordnungen – untersuchen.

Felix Hausdorff über die Ordinalzahlen

„Der Leser wird an dieser Stelle gern einen kurzen Blick rückwärts tun und der genialen Schöpfung G. Cantors, dem System der Ordnungszahlen, seine Bewunderung nicht versagen. Die letzten Betrachtungen über den Anfang dieses Systems zeigen, daß die paradox erscheinende Idee, über die endliche Zahlenreihe hinaus den Zählprozeß fortzusetzen, wirklich ausführbar ist, und zwar nicht in einer nebelhaften Weise mit fragwürdigen Unendlichkeitssymbolen wie ∞ , sondern nach einem präzisen Gesetz, das an jeder Stelle des Zahlensystems die nunmehr folgende Zahl als Typus der Menge aller vorangehenden Zahlen eindeutig bestimmt. Für die wohlgeordneten Mengen ist damit auch das Postulat erfüllt, daß der Zählprozeß nicht nur die ganze Menge, sondern auch ihre Elemente ‚zählen‘, ihnen nämlich bestimmte Zahlzeichen als Nummern oder Indices zuordnen sollte; denn die Elemente einer wohlgeordneten Menge A vom Typus α werden eben durch die Zahlen der Menge $W(\alpha)$ in diesem Sinne gezählt, d. h. umkehrbar eindeutig repräsentiert.“

(Felix Hausdorff 1914, „Grundzüge der Mengenlehre“)

7. Transfinite Induktion und Rekursion

Die transfinite Induktion und Rekursion ermöglicht Beweise und Definitionen *entlang* der Ordinalzahlen. Wir betrachten zunächst das Schema der vollständigen Induktion für die natürlichen Zahlen, das wir schon an verschiedenen Stellen verwendet haben, etwas genauer. Eine Form dieses Prinzips lautet:

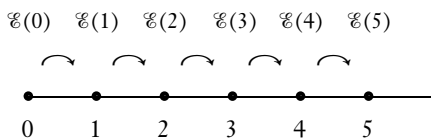
Sei $\mathcal{E}(n)$ eine Aussage. Es gelte:

(1) $\mathcal{E}(0)$

(2) Für alle $n \in \mathbb{N}$ gilt:

$\mathcal{E}(n)$ folgt $\mathcal{E}(n+1)$.

Dann gilt $\mathcal{E}(n)$ für alle $n \in \mathbb{N}$.



Im Gegensatz zur reinen Zahlentheorie, wo man dieses Prinzip als Axiom annimmt, ist in der Mengenlehre die Induktion ein beweisbarer Satz:

Annahme, es gibt ein $n \in \mathbb{N}$ mit $\text{non } \mathcal{E}(n)$. Dann ist $A = \{n \in \mathbb{N} \mid \text{non } \mathcal{E}(n)\} \neq \emptyset$. Da $\langle \mathbb{N}, < \rangle$ wohlgeordnet ist, hat also A ein kleinstes Element n . Nach (1) ist $n \neq 0$. Sei also $n = m + 1$. Nach Definition von A und n gilt $\mathcal{E}(m)$. Nach (2) folgt aus $\mathcal{E}(m)$ aber $\mathcal{E}(m+1)$, also $\mathcal{E}(n)$, *Widerspruch!*

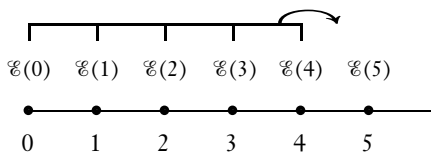
Entscheidend ist also diese Eigenschaft: Existiert ein Gegenbeispiel zur Behauptung, so existiert ein kleinstes Gegenbeispiel. Dies ist aber gerade die Wohlordnungseigenschaft, und wir können nun mit einem ähnlichen Argument das Prinzip der Induktion auf die transfiniten Zahlen ausdehnen. Wir beweisen die transfinite Induktion in einer starken und kompakten Form, die folgender Version der vollständigen Induktion für die natürlichen Zahlen entspricht:

Sei $\mathcal{E}(n)$ eine Aussage.

Für alle $n \in \mathbb{N}$ gelte:

(+) Gilt $\mathcal{E}(m)$ für alle $m < n$,
so gilt $\mathcal{E}(n)$.

Dann gilt $\mathcal{E}(n)$ für alle $n \in \mathbb{N}$.



Der Beweis verläuft wie eben:

Ein kleinstes Gegenbeispiel liefert einen Widerspruch zur Voraussetzung (+). (Aus (+) folgt $\mathcal{E}(0)$. Denn für $n = 0$ ist „für alle $m < n$ gilt $\mathcal{E}(m)$ “ sicher richtig. Also gilt $\mathcal{E}(n)$ nach (+).)

Die Induktion wird angewendet, um „ $\mathcal{E}(n)$ gilt für alle $n \in \mathbb{N}$ “ zu beweisen. Hierzu kann man (1) und (2) zeigen, oder man zeigt (+). Letzteres ist die stärkere Form der Induktion, da man für den Beweis von $\mathcal{E}(n)$ die Gültigkeit der Aus-

sage $\mathcal{E}$ für *alle* Vorgänger von n als bereits bewiesen annehmen darf, nicht nur die Gültigkeit für einen unmittelbaren Vorgänger wie in (2).

Im folgenden bezeichnen wir mit $\mathcal{E}(\alpha)$ eine Aussage über Ordinalzahlen, die auch Parameter enthalten darf. Z. B. ist ω ein Parameter in $\mathcal{E}(\alpha)$ = „für alle Ordinalzahlen $\alpha \geq \omega$ gilt $|\alpha| = |\alpha + 1|$ “.

Satz (*Beweis durch transfinite Induktion*)

Sei $\mathcal{E}(\alpha)$ eine Aussage. Für alle Ordinalzahlen α gelte:

(+) Gilt $\mathcal{E}(\beta)$ für alle $\beta < \alpha$, so gilt $\mathcal{E}(\alpha)$.

Dann gilt $\mathcal{E}(\alpha)$ für alle Ordinalzahlen α .

Beweis

Annahme nicht. Dann existiert eine Ordinalzahl γ , für die $\mathcal{E}(\gamma)$ falsch ist.

Dann ist also

$$A = \{ \alpha \leq \gamma \mid \mathcal{E}(\alpha) \text{ ist falsch} \}$$

eine nichtleere Menge von Ordinalzahlen. Also besitzt A ein kleinstes Element α . Wegen Minimalität von α gilt dann $\mathcal{E}(\beta)$ für alle $\beta < \alpha$.

– Nach (+) gilt also $\mathcal{E}(\alpha)$, im Widerspruch zu $\alpha \in A$.

Es genügt also, (+) zu zeigen, um eine Aussage $\mathcal{E}(\alpha)$ für alle Ordinalzahlen α nachzuweisen. Hierfür muß man zwar auch noch für ein beliebiges α die Aussage $\mathcal{E}(\alpha)$ beweisen, darf dabei aber annehmen, daß $\mathcal{E}(\beta)$ für *alle* $\beta < \alpha$ bereits bewiesen ist. Man hat also zur Unterstützung eine – oft stillschweigend gemachte – *Induktionsvoraussetzung* zur Verfügung.

In der Praxis wird häufig die folgende dreiteilige Version der transfiniten Induktion verwendet, da man zum Nachweis von $\mathcal{E}(\alpha)$ nicht immer die Gültigkeit der Aussage für alle Vorgänger von α benötigt:

Ziel: $\mathcal{E}(\alpha)$ gilt für alle Ordinalzahlen α .

Zeige dies in drei Teilen:

- (1) Es gilt $\mathcal{E}(0)$. (*Induktionsanfang* $\alpha = 0$)
- (2) Aus $\mathcal{E}(\alpha)$ folgt $\mathcal{E}(\alpha + 1)$
für alle Ordinalzahlen α . (*Nachfolgerschritt* von α nach $\alpha + 1$)
- (3) Aus $\mathcal{E}(\alpha)$ für alle $\alpha < \lambda$ folgt $\mathcal{E}(\lambda)$
für alle Limesordinalzahlen λ . (*Limeschritt* λ)

Man sieht sofort, daß die Argumentation über einen „kleinsten Ausreißer“ auch unter diesen Voraussetzungen zu einem Widerspruch führt, falls $\mathcal{E}(\gamma)$ für eine Ordinalzahl γ falsch wäre. Denn ein kleinster Ausreißer ist entweder 0, ein Nachfolger $\alpha + 1$ oder ein Limes λ . Und aus (1), (2) oder (3) folgt dann im jeweiligen Fall der Widerspruch.

In dieser Form ist die transfinite Induktion soweit als möglich analog zur üblichen „für 0 und von n nach $n + 1$ “-Induktion der natürlichen Zahlen formuliert. Lediglich ein Limeschritt λ kommt hinzu, der zeigt, daß die Richtigkeit der Aussage nicht an einer Limesstufe zum ersten Mal verloren geht.

Ob man (+) wie im Satz oben oder die dreiteilige Form (oder andere Varianten) nachweist, hängt von der Hartnäckigkeit der Aussage $\mathcal{E}(\alpha)$ ab. Die Unterscheidung in Nachfolger- und Limesordinalzahlen ist aber den Ordinalzahlen so eigen-tümlich, daß die dreiteilige Form bei weitem am häufigsten verwendet wird.

Hausdorff (1914): „Hier muß zunächst auf eine weitgehende Analogie mit der Reihe der endlichen Zahlen ... hingewiesen werden, nämlich auf die Anwendbarkeit des sogenannten vollständigen Induktionsschlusses. Der Leser kennt aus zahllosen Beispielen den Schluß von v auf $v + 1$, der besagt: eine Aussage $f(v)$ bezüglich der endlichen Zahl v ist für jedes v richtig, falls $f(0)$ richtig ist und falls aus der Richtigkeit von $f(v)$ auf die von $f(v + 1)$ geschlossen werden kann. Im Gebiete der endlichen und unendlichen Ordnungszahlen gilt nun ein ähnliches Schlußverfahren, nämlich:

Eine Aussage $f(\alpha)$ bezüglich der Ordnungszahl α ist für jedes α richtig, sobald $f(0)$ richtig ist und sobald aus der Richtigkeit aller $f(\xi)$ für $\xi < \alpha$ auf die Richtigkeit von $f(\alpha)$ geschlossen werden kann.“

Viele weitere Varianten werden oft ohne Kommentar verwendet, etwa die transfinite Induktion innerhalb eines Intervalls: Sind α_0, α_1 Ordinalzahlen und ist $\alpha_0 < \alpha_1$, so kann man „ $\mathcal{E}(\alpha)$ gilt für alle α mit $\alpha_0 \leq \alpha < \alpha_1$ “ beweisen, indem man zeigt:

- (1) $\mathcal{E}(\alpha_0)$,
- (2) $\mathcal{E}(\alpha)$ folgt $\mathcal{E}(\alpha + 1)$ für alle $\alpha \geq \alpha_0$ mit $\alpha + 1 < \alpha_1$,
- (3) $\mathcal{E}(\alpha)$ für alle α mit $\alpha_0 \leq \alpha < \lambda$ folgt $\mathcal{E}(\lambda)$ für alle Limesordinalzahlen $\lambda < \alpha_1$.

Transfinite Rekursion

Neben dem Beweis durch transfinite Induktion gibt es die Definition durch transfinite Rekursion, die manchmal auch induktive Definition genannt wird. Hierbei wollen wir der Reihe nach für jede Ordinalzahl α ein Objekt $\mathcal{G}(\alpha)$ in einer bestimmten Weise konstruieren, und dabei auf alle bereits konstruierten Objekte $\mathcal{G}(\beta)$ für $\beta < \alpha$ zurückgreifen. Die „bestimmte Weise“ ist gegeben durch eine Operation auf dem Universum V aller Objekte:

Definition (Eigenschaften als Operationen)

Eine Eigenschaft $\mathcal{F}(x, y)$ heißt eine *Operation auf V* , falls für jedes Objekt x genau ein Objekt y existiert mit $\mathcal{F}(x, y)$. Analog heißt $\mathcal{G}(\alpha, y)$ eine *Operation auf den Ordinalzahlen*, falls für jede Ordinalzahl α genau ein Objekt y existiert mit $\mathcal{G}(\alpha, y)$.

Wir schreiben auch $\mathcal{F}(x) = y$, falls $\mathcal{F}(x, y)$ gilt.

Daß wir hier von Eigenschaften reden anstatt einfach von Funktionen hängt damit zusammen, daß die Vielheit V aller Objekte keine Menge ist und also auch eine „Funktion“ $\mathcal{F} : V \rightarrow V$ keine Menge sein kann. Beispiele für Operationen auf V sind:

$$\begin{aligned}\mathcal{F}(x) &= \bigcup x \quad (\text{d.h. } \mathcal{F}(x, y) \text{ gilt genau dann, wenn } y = \bigcup x), \\ \mathcal{F}(x) &= \mathcal{P}(x), \quad \text{jeweils mit der Konvention } \bigcup x = \mathcal{P}(x) = x, \text{ falls } x \text{ Grundobjekt,} \\ \mathcal{F}(x) &= \begin{cases} x \cup \mathbb{N}, & \text{falls } x \text{ endliche Menge,} \\ x, & \text{falls } x \text{ unendliche Menge oder Grundobjekt.} \end{cases}\end{aligned}$$

Hierbei ist x ein beliebiges Objekt des Universums. Im letzten Beispiel ist $\mathbb{N}$ ein Parameter der Operation.

Bevor wir den Rekursionssatz formulieren und beweisen, werfen wir wieder einen Blick auf die natürlichen Zahlen. Ein Beispiel für eine rekursive Definition einer Operation $\mathcal{G}$ auf $\mathbb{N}$ – hier einfach eine Funktion auf $\mathbb{N}$ – ist die Fakultät:

- (1) $0! = 1,$ (Rekursionsanfang)
- (2) $(n+1)! = n! \cdot (n+1)$ für $n \in \mathbb{N}.$ (Rekursionsschritt von n nach $n+1$)

Für die Fakultät $\mathcal{G}: \mathbb{N} \rightarrow \mathbb{N}$, $\mathcal{G}(n) = n!$ für $n \in \mathbb{N}$, gilt also die Rekursionsgleichung

$$\mathcal{G}(n+1) = \mathcal{G}(n) \cdot (n+1) = \mathcal{F}(\mathcal{G}(n), n),$$

mit der Operation $\mathcal{F}(m, n) = m \cdot (n+1)$.

Ein rekursiver Funktionswert $\mathcal{G}(n)$ kann aber auch von mehreren Werten $\mathcal{G}(m)$, $m < n$, abhängen. Ein bekanntes Beispiel hierfür liefern die Fibonacci-Zahlen 1, 1, 2, 3, 5, ... Die n -te Fibonacci-Zahl $\mathcal{G}(n)$ ist hierbei rekursiv definiert durch $\mathcal{G}(0) = \mathcal{G}(1) = 1$, $\mathcal{G}(n) = \mathcal{G}(n-1) + \mathcal{G}(n-2)$ für $n \geq 2$.

Allgemein ist der Rückgriff auf die ganze Liste der bereits definierten Funktionswerte möglich, und die Rekursionsgleichung lautet dann:

$$\mathcal{G}(n) = \mathcal{F}(\langle \mathcal{G}(n') \mid n' < n \rangle) = \mathcal{F}(\mathcal{G} \upharpoonright \bar{n}),$$

wobei $\mathcal{F}$ eine gegebene Operation ist, die die „bestimmte Weise“ beschreibt, wie $\mathcal{G}(n)$ aus $\mathcal{G}(0), \dots, \mathcal{G}(n-1)$ hervorgeht. Speziell ist $\mathcal{G}(0) = \mathcal{F}(\emptyset)$. (Der häufig zur Definition von $\mathcal{G}(n) = \mathcal{F}(\mathcal{G} \upharpoonright \bar{n})$ benötigte Index n selbst läßt sich aus $\mathcal{G} \upharpoonright \bar{n}$ zurückgewinnen, denn es ist $\bar{n} = \text{dom}(\mathcal{G} \upharpoonright \bar{n})$.)

Für Ordinalzahlen lautet nun der allgemeine Rekursionssatz:

Satz (transfiniten Rekursionssatz)

Sei $\mathcal{F}$ eine Operation auf V . Dann existiert genau eine Operation $\mathcal{G}$ auf den Ordinalzahlen, sodaß gilt:

$$(+)\quad \text{Für alle Ordinalzahlen } \alpha \text{ ist } \mathcal{G}(\alpha) = \mathcal{F}(\mathcal{G} \upharpoonright W(\alpha)),$$

wobei $\mathcal{G} \upharpoonright W(\alpha)$ die Funktion $g_\alpha: W(\alpha) \rightarrow V$ ist mit $g_\alpha(\beta) = \mathcal{G}(\beta)$ für alle $\beta < \alpha$.

Die *Rekursionsgleichung* (+) können wir auch als $\mathcal{G}(\alpha) = \mathcal{F}(\langle \mathcal{G}(\beta) \mid \beta < \alpha \rangle)$ schreiben. So kommt besonders deutlich zum Ausdruck, daß das im α -ten Schritt mittels $\mathcal{F}$ rekursiv definierte Objekt $\mathcal{G}(\alpha)$ von allen in früheren Schritten $\beta < \alpha$ konstruierten Objekten $\mathcal{G}(\beta)$ abhängen kann.

Beweis

Wir zeigen zunächst die Existenz von $\mathcal{G}$.

Wir nennen für eine Ordinalzahl α eine Funktion $g : W(\alpha) \rightarrow V$ eine α -Rekursion gemäß $\mathcal{F}$, falls gilt:

$$g(\beta) = \mathcal{F}(g \upharpoonright W(\beta)) \text{ für alle } \beta < \alpha.$$

(Derartige Funktionen g sind die Anfangsstücke der gesuchten Operation $\mathcal{G}$.)

Es gilt nun:

- (i) Für alle Ordinalzahlen α gilt:
Sind g und h α -Rekursionen gemäß $\mathcal{F}$, so ist $g = h$.
- (ii) Für alle Ordinalzahlen α existiert eine α -Rekursion gemäß $\mathcal{F}$.

Beweis hierzu

zu (i): Sei α eine Ordinalzahl und seien g, h α -Rekursionen gemäß $\mathcal{F}$.

Wir zeigen $g(\beta) = h(\beta)$ durch Induktion nach $\beta < \alpha$.

Induktionsschritt $\beta < \alpha$: Es gelte $g(\beta') = h(\beta')$ für alle $\beta' < \beta$.

Dann ist also $g \upharpoonright W(\beta) = h \upharpoonright W(\beta)$, also

$$g(\beta) = \mathcal{F}(g \upharpoonright W(\beta)) = \mathcal{F}(h \upharpoonright W(\beta)) = h(\beta).$$

zu (ii): Beweis durch Induktion nach α .

Induktionsanfang $\alpha = 0$:

Die leere Menge ist eine 0-Rekursion gemäß $\mathcal{F}$.

Induktionsschritt von α nach $\alpha + 1$:

Nach Induktionsvoraussetzung

existiert eine α -Rekursion g gemäß $\mathcal{F}$.

Sei $g' = g \cup \{(\alpha, \mathcal{F}(g))\}$.

Dann ist g' eine $(\alpha + 1)$ -Rekursion gemäß $\mathcal{F}$.

Limeschritt λ :

Nach Induktionsvoraussetzung und (i) existiert

für jedes $\alpha < \lambda$ eine eindeutige α -Rekursion g_α gemäß $\mathcal{F}$.

Wir setzen:

$$g = \bigcup_{\alpha < \lambda} g_\alpha.$$

Dann ist g eine λ -Rekursion gemäß $\mathcal{F}$.

Wir definieren nun eine Operation $\mathcal{G}$ auf den Ordinalzahlen durch:

$\mathcal{G}(\alpha) =$ „das x mit $g(\alpha) = x$, wobei g die eindeutige $(\alpha + 1)$ -Rekursion gemäß $\mathcal{F}$ ist“.

Nach Konstruktion gilt dann $\mathcal{G}(\alpha) = \mathcal{F}(\mathcal{G} \upharpoonright W(\alpha))$ für alle Ordinalzahlen α . Dies zeigt die Existenz.

Sind nun $\mathcal{G}_1, \mathcal{G}_2$ Operationen mit (+) wie im Rekursionssatz, so sind $\mathcal{G}_1 \upharpoonright W(\alpha)$ und $\mathcal{G}_2 \upharpoonright W(\alpha)$ α -Rekursionen gemäß $\mathcal{F}$ für alle Ordinalzahlen α . Also gilt $\mathcal{G}_1 \upharpoonright W(\alpha) = \mathcal{G}_2 \upharpoonright W(\alpha)$ für alle α nach (i).

– Dann gilt aber $\mathcal{G}_1(\alpha) = \mathcal{G}_2(\alpha)$ für alle Ordinalzahlen α .

Hausdorff (1914): „Wir können aber nicht nur induktiv schließen, sondern auch induktiv definieren. Bedeutet $f(\alpha)$ jetzt nicht eine Aussage hinsichtlich α , sondern ein der Zahl α zugeordnetes Ding, eine Funktion von $\alpha \dots$, so lautet das induktive Definitionsverfahren:

$f(\alpha)$ ist für jedes α definiert, sobald $f(0)$ definiert ist und sobald vermöge der Definition aller $f(\xi)$ für $\xi < \alpha$ auch $f(\alpha)$ definiert ist.“

Hausdorff hat diese Möglichkeit der rekursiven Definition für so selbstverständlich erachtet, daß er keine weitere Begründung angibt.

Beispiele für transfinite Rekursionen

Wir behandeln zunächst die vier „transfiniten Prozesse“ aus den beiden ersten Kapiteln dieses Abschnitts.

(1) Der Algorithmus des Abtragens für beliebige Mengen

Sei M eine Menge. Wir tragen M entlang der Ordinalzahlen rekursiv ab (vgl. hierzu auch den Beweis des Wohlordnungssatzes):

Sei $f: \mathcal{P}(M) - \{\emptyset\} \rightarrow M$ eine Funktion mit $f(A) \in A$ für alle nichtleeren $A \subseteq M$. Wir definieren solange möglich $x_\alpha \in M$ rekursiv durch:

$$x_\alpha = f(M - \{x_\beta \mid \beta < \alpha\}).$$

Die Definition bricht an einer Stelle $\alpha^* < |M|^+$ ab, da wir sonst eine Injektion von $W(|M|^+)$ nach M erhalten würden. Für α^* ist dann aber $M - \{x_\alpha \mid \alpha < \alpha^*\} \notin \text{dom}(f)$, also leer. Somit ist $\langle x_\alpha \mid \alpha < \alpha^* \rangle$ eine Aufzählung von M , gewonnen durch „Abtragen von M entlang der Ordinalzahlen gemäß f .“ Die Funktion f wählt an jeder Stelle ein Element aus dem Resthaufen.

Der verwendete Ausdruck „definiere x_α solange möglich“ ist eine bequeme Sprechweise. Um ihn zu vermeiden, kann man so vorgehen: Sei x^* ein Objekt mit $x^* \notin M$. Man definiert dann x_α für alle Ordinalzahlen α durch:

$$x_\alpha = f(M - \{x_\beta \mid \beta < \alpha\}), \text{ falls } M - \{x_\beta \mid \beta < \alpha\} \neq \emptyset, \quad x_\alpha = x^* \text{ sonst.}$$

Weiter setzt man $\alpha^* =$ „das kleinste α mit $x_\alpha = x^*$ “, und hat dann die gleiche Aufzählung $\langle x_\alpha \mid \alpha < \alpha^* \rangle$ wie zuvor.

Eine zweite Bemerkung: Warum schreiben wir den Rekursionsschritt nicht einfach in der Form: $x_\alpha =$ „ein $x \in M - \{x_\beta \mid \beta < \alpha\}$ “? Der Grund ist, daß wir für den Rekursionssatz eine feste Operation $\mathcal{F}$ verwenden müssen, die in unserem Fall etwa lautet:

$$\mathcal{F}(G) = f(M - \text{rng}(G)), \text{ falls } G \text{ Funktion, und } \mathcal{F}(G) = \emptyset \text{ sonst.}$$

Die Aufzählung $\langle x_\alpha \mid \alpha < \alpha^* \rangle$, d.h. die Funktion $g: W(\alpha^*) \rightarrow M$ mit $g(\alpha) = x_\alpha$ für $\alpha < \alpha^*$, induziert eine Wohlordnung auf M . Warum hat Cantor nicht auf diese Weise den Wohlordnungssatz bewiesen? Nicht, weil er die transfinite Rekursion nicht kannte. Und die Idee dieses Beweises hatte er klar vor Augen. Das Problem ist: Man muß zeigen, daß man fertig wird, d.h. daß es eine Stelle α^* gibt, an

der die Rekursion abbricht. Wir haben oben argumentiert, daß wir andernfalls für $\kappa = |M|$ eine Injektion von $W(\kappa^+)$ nach $W(\kappa)$ erhalten würden. Hierzu verwenden wir implizit, daß sich M bijektiv auf ein Anfangsstück der Ordinalzahlen abbilden läßt, d. h. daß wir $|M|$ als Ordinalzahl auffassen können. Und hierzu wird der Wohlordnungssatz bereits benutzt!

Wir geben nach der Diskussion weiterer Beispiele ein Argument für das Fertigwerden, das auf dem Satz von Hartogs ruht, und das einen neuen Beweis für den Wohlordnungssatz liefert. Weiter besprechen wir auch ein Argument von Cantor, mit dem er das Fertigwerden begründete.

Durch paralleles rekursives Abtragen zweier Mengen läßt sich nun leicht der Vergleichbarkeitssatz beweisen, und wir haben damit einen Beweis gefunden, der der intuitiven Begründung für die Gültigkeit des Satzes aus dem ersten Abschnitt genau entspricht:

Übung

Seien M, N Mengen.

Beweisen Sie „ $|M| \leq |N|$ oder $|N| \leq |M|$ “ durch „rekursives Abtragen.“

[Seien $\mathcal{P}(M)^- = \mathcal{P}(M) - \{\emptyset\}$, $\mathcal{P}(N)^- = \mathcal{P}(N) - \{\emptyset\}$.

Weiter sei $f : \mathcal{P}(M)^- \times \mathcal{P}(N)^- \rightarrow M \times N$ eine Funktion mit $f(A, B) \in A \times B$.

Wir definieren rekursiv solange möglich:

$(x_\alpha, y_\alpha) = f(M - \{x_\beta \mid \beta < \alpha\}, N - \{y_\beta \mid \beta < \alpha\})$.

Es existiert ein α^* , an dem die Rekursion abbricht ($\alpha^* < \min(|N|^+, |M|^+)$).

Dann zeigt g oder g^{-1} die Behauptung mit $g = \{(x_\alpha, y_\alpha) \mid \alpha < \alpha^*\}$.]

Weiter lassen sich die Maximalprinzipien mit Hilfe der transfiniten Rekursion sehr einfach beweisen:

Übung

Beweisen Sie den Satz von Zermelo-Zorn (oder das Hausdorffsche Maximalprinzip) mit Hilfe transfiniten Rekursion.

[Wir konstruieren rekursiv eine strikt aufsteigende Folge $\langle x_\alpha \mid \alpha \leq \alpha^* \rangle$ von Elementen der partiellen Ordnung, bis wir zu einem maximalen Element x_{α^*} gelangen. Im Limeschritt wird die Voraussetzung der Existenz oberer Schranken für linear geordnete Teilmengen benutzt.]

(2) Mengen immer größerer Mächtigkeit durch iterierte Potenzmengenbildung

Sei M eine Menge. Dann sind die α -ten Potenzen $\mathcal{P}^\alpha(M)$ von M rekursiv definiert durch:

$$\mathcal{P}^0(M) = M,$$

$$\mathcal{P}^{\alpha+1}(M) = \mathcal{P}(\mathcal{P}^\alpha(M)) \quad \text{für } \alpha \text{ Ordinalzahl,}$$

$$\mathcal{P}^\lambda(M) = \bigcup_{\alpha < \lambda} \mathcal{P}^\alpha(M) \quad \text{für } \lambda \text{ Limes.}$$

Die Kardinalzahlen $\beth_\alpha$ [$\beth = \text{beth}$] sind definiert durch:

$\beth_\alpha = |\mathcal{P}^\alpha(\mathbb{N})|$ für α Ordinalzahl.

Sie können auch rekursiv definiert werden durch:

$$\beth_0 = \omega,$$

$$\beth_{\alpha+1} = 2^{\beth_\alpha} \quad \text{für } \alpha \text{ Ordinalzahl,}$$

$$\beth_\lambda = \sup_{\alpha < \lambda} \beth_\alpha \quad \text{für } \lambda \text{ Limes.}$$

(3) Schichtung von V

Die *Neumann-Zermelo-Stufen* V_α sind rekursiv definiert durch

$$V_0 = \{x \mid x \text{ ist Grundobjekt}\},$$

$$V_{\alpha+1} = \mathcal{P}(V_\alpha) \cup V_0 \quad \text{für } \alpha \text{ Ordinalzahl,}$$

$$V_\lambda = \bigcup_{\alpha < \lambda} V_\alpha \quad \text{für } \lambda \text{ Limes.}$$

Man zeigt leicht durch Induktion nach α , daß jede Menge $x \in V_\alpha$ auch eine Teilmenge von V_α ist. Folglich gilt $V_\alpha \subseteq V_\beta$ für alle Ordinalzahlen $\alpha \leq \beta$ (durch Induktion nach β bei festem α). Die V_α -Stufen sind also transitiv und $\subseteq$ -aufsteigend. Man nennt die Folge der V_α -Stufen deswegen auch eine *kumulative Hierarchie*.

Es gilt $|V_{n+1}| = 2^n$ für alle $n \in \omega$ und $|V_{\omega+\alpha}| = \beth_\alpha$ für alle α , falls wir annehmen, daß es keine Grundobjekte gibt, wie es heute in der axiomatischen Mengenlehre üblicherweise gemacht wird (vgl. Abschnitt 3). Unter dieser Annahme gilt dann $V_{\alpha+1} = \mathcal{P}(V_\alpha)$ für alle α und damit $V_\alpha = \mathcal{P}^\alpha(\emptyset)$ für alle α .

Der folgende Satz charakterisiert diejenigen Objekte, die in der V_α -Hierarchie eingefangen werden. Das Kriterium betrifft unendlich absteigende $\in$ -Ketten, die uns bei der Mirimanov-Paradoxie bereits begegnet sind.

Satz (*V und die Neumann-Zermelo-Hierarchie*)

Sei x ein Objekt. Dann sind äquivalent:

- (i) Es gibt eine Ordinalzahl α mit $x \in V_\alpha$.
- (ii) Es gibt keine Mengen $x_0 \ni x_1 \ni \dots \ni x_n \ni \dots$, $n \in \mathbb{N}$, mit $x_0 = x$.

Beweis

(i) $\hookrightarrow$ (ii): Wir zeigen durch Induktion nach α :

Ist $x \in V_\alpha$, so gibt es keine Mengen $x_0 \ni \dots \ni x_n \ni \dots$, $n \in \mathbb{N}$, mit $x_0 = x$.

Induktionsanfang $\alpha = 0$:

Ist $x \in V_0$, so ist x Grundobjekt und die Aussage ist trivial.

Induktionsschritt α nach $\alpha + 1$: Sei $x \in V_{\alpha+1}$.

Annahme, es gibt $x_0 \ni x_1 \ni \dots \ni x_n \ni \dots$, $n \in \mathbb{N}$, mit $x_0 = x$.

Dann ist $x \notin V_0$, also $x \subseteq V_\alpha$. Damit ist $x_1 \in V_\alpha$ und es gilt $x_1 \ni x_2 \ni x_3 \ni \dots$, im *Widerspruch* zur Induktionsvoraussetzung.

Limesschritt λ : Sei $x \in V_\lambda$.

Annahme, es gibt $x_0 \ni x_1 \ni \dots \ni x_n \ni \dots$, $n \in \mathbb{N}$, mit $x_0 = x$.

Wegen $V_\lambda = \bigcup_{\alpha < \lambda} V_\alpha$ ist dann $x \in V_\alpha$ für ein $\alpha < \lambda$, und wie eben ergibt sich ein *Widerspruch* zur Induktionsvoraussetzung.

(ii) $\curvearrowright$ (i):

Wir zeigen non (i) $\curvearrowright$ non (ii). Sei also $x \notin V_\alpha$ für alle α .

Dann ist x kein Grundobjekt, da sonst $x \in V_0$.

Wir definieren rekursiv Mengen x_n für $n \in \mathbb{N}$ durch $x_0 = x$ und :

$x_{n+1} =$ „ein $y \in x_n$ mit $y \notin V_\alpha$ für alle Ordinalzahlen α “.

Ein solches y existiert (und ist dann kein Grundobjekt):

Annahme nicht. Für alle $y \in x_n$ sei dann

$\alpha_y =$ „die kleinste Ordinalzahl α mit $y \in V_\alpha$ “,

und sei $A = \{ \alpha_y \mid y \in x_n \}$. Weiter sei $\beta = \sup(A)$.

Dann gilt $x_n \subseteq V_\beta$, also $x_n \in V_{\beta+1}$, *Widerspruch!*

– Dann ist aber $x_0 \ni x_1 \ni x_2 \ni \dots$, also gilt non(ii) für x .

Unter der Annahme:

(F) Es gibt keine Mengen $x_0 \ni x_1 \ni x_2 \ni \dots \ni x_n \ni \dots$, $n \in \mathbb{N}$.

wird also jedes Objekt in einem V_α eingefangen. In jedem Falle aber enthält die V_α -Hierarchie sehr viele Objekte (vgl. auch das Mirimanovsche Paradoxon). Die Hierarchie läßt genau die schwarzen Löcher $x_0 \ni x_1 \ni x_2 \ni \dots$ aus.

In der axiomatischen Mengenlehre wird die Gültigkeit dieser restriktiven Aussage (F) – aus der sofort $x \notin x$ für alle Mengen x folgt – durch das sog. Fundierungsaxiom garantiert. Die Rechtfertigung ist eine Verstärkung des iterativen Mengenkonzept: Mengenbildungen sind nicht nur iterativ, sondern alles wird aus den Grundobjekten oder aus dem Nichts $\emptyset$ durch iterierte Mengenbildung erzeugt. So ergibt sich ein Atlas des Mengenuniversums: Unter (F) liegt das Universum in der Neumann-Zermelo-Hierarchie ganz und strukturiert vor uns ausgebreitet, und die Ordinalzahlen bilden einen hellen Pfad, der alle Fernen erreicht.

Übung

Für alle Neumann-Zermelo-Ordinalzahlen α gilt:

$\alpha =$ „die kleinste Ordinalzahl β mit $\alpha \subseteq V_\beta$ “.

Wir diskutieren noch zwei wichtige Konsequenzen von (F): Ränge und Trunkierungen. Unter (F) können wir einen natürlichen Rang für alle Objekte definieren. Wir setzen für jedes Objekt x :

$\mathcal{R}(x) =$ „das kleinste α mit $x \in V_{\alpha+1}$ “.

Nach obiger Übung ist $\mathcal{R}(\alpha) = \alpha$ für alle Ordinalzahlen α .

Durch eine derartige Rangdefinition erhalten wir eine Schichtung des Universums wie in den Diagrammen des ersten Abschnitts (1.13). Wir setzen:

$Z_\alpha = \{ x \mid \mathcal{R}(x) = \alpha \}$ für α Ordinalzahl.

Weiter können wir zu große Zusammenfassungen zu Mengen trunkieren: Sei $\mathcal{E}(x)$ eine Eigenschaft (mit Parametern). Wir setzen (Trick von Dana Scott):

$S_{\mathcal{E}(x)} = \{ x \mid \mathcal{E}(x) \text{ und für alle } y \text{ mit } \mathcal{E}(y) \text{ gilt } \mathcal{R}(x) \leq \mathcal{R}(y) \}$.

Dann gilt $S_{\mathcal{E}(x)} \subseteq V_\alpha$ für alle α , für die ein $x \in V_\alpha$ existiert mit $\mathcal{E}(x)$. Wir sammeln also alle x mit $\mathcal{E}(x)$ in derjenigen Stufe auf, in der zum ersten Mal überhaupt x auftauchen mit $\mathcal{E}(x)$. So entsteht eine Teilmenge einer Stufe, also eine Menge. In dieser Weise können wir dann z. B. die Mächtigkeit $|M|$ von M definieren als die Menge $S_{\mathcal{E}(N)}$ mit

$$\mathcal{E}(N) = „|N| = |M|“$$

mit M als Parameter. Oder den Ordnungstyp o.t. $\langle M, < \rangle$ als $S_{\mathcal{E}(\langle N, < \rangle)}$ mit

$$\mathcal{E}(\langle N, < \rangle) = „\langle M, < \rangle \text{ und } \langle N, < \rangle \text{ sind isomorph.}“$$

Das alles ist nicht gerade schön, aber möglich – unter Annahme von (F).

(4) Ableitung einer Punktmenge

Sei $P \subseteq \mathbb{R}$. Dann sind die *iterierten Ableitungen* $P^{(\alpha)}$ von P rekursiv definiert durch:

$$\begin{aligned} P^{(0)} &= P, \\ P^{(\alpha+1)} &= P^{(\alpha)'} && \text{für } \alpha \text{ Ordinalzahl,} \\ P^{(\lambda)} &= \bigcap_{\alpha < \lambda} P^{(\alpha)} && \text{für } \lambda \text{ Limes,} \end{aligned}$$

wobei für $Q \subseteq \mathbb{R}$ die Menge Q' aus den Häufungspunkten von Q besteht.

Ein weiteres Beispiel für eine transfinite Rekursion liefern die Cantorsche Alephs, die die unendlichen Kardinalzahlen aufzählen.

(5) Die Aleph-Folge der transfiniten Kardinalzahlen

Die *unendlichen Kardinalzahlen* $\aleph_\alpha$ [Aleph α] sind rekursiv definiert durch:

$$\begin{aligned} \aleph_0 &= \omega, \\ \aleph_{\alpha+1} &= (\aleph_\alpha)^+ && \text{für } \alpha \text{ Ordinalzahl,} \\ \aleph_\lambda &= \sup_{\alpha < \lambda} \aleph_\alpha && \text{für } \lambda \text{ Limes.} \end{aligned}$$

Neben der Bezeichnung $\aleph_\alpha$ ist auch ω_α gebräuchlich.

Die Folge der Alephs ist die monotone Aufzählung aller unendlichen Kardinalzahlen. Die Kontinuumshypothese besagt $\aleph_1 = \mathfrak{c}_1$. Die allgemeine Kontinuumshypothese kann man schreiben als $\aleph_\alpha = \mathfrak{c}_\alpha$ für alle Ordinalzahlen α .

Transfinite Induktion und Rekursion gehen auf Cantor zurück. Die typische dreiteilige Form obiger Beispiele für rekursive Definitionen, und ein induktiver Beweis, daß dadurch eindeutig bestimmte Objekte definiert werden, tauchen in klarer Weise zum ersten Mal bei der Definition der Potenzierung von Ordinalzahlen auf (Cantor 1897, § 18, vgl. das nächste Kapitel). Innerhalb eines formalen – oder genauer: formalisierbaren – Rahmens hat die transfinite Rekursion dann von Neumann behandelt (1923).

Nicht alle Rekursionen laufen über alle Ordinalzahlen. Ein fundamentales Beispiel für eine Rekursion bis hinauf zu ω_1 ist die Borel-Hierarchie, benannt nach Émile Borel (1871 – 1956).

(6) Die Borel-Hierarchie

Abgeschlossene Teilmengen von $\mathbb{R}$ sind stabil unter abzählbaren Schnitten, dagegen führen abzählbare Vereinigungen von abgeschlossenen Mengen aus den abgeschlossenen Mengen hinaus. Analoges gilt für die offenen Mengen. Die Generationsprinzipien „abzählbare Vereinigung“ und „abzählbarer Durchschnitt“ führen nun durch iterierte Anwendung zu immer komplizierteren Teilmengen von $\mathbb{R}$. Hierzu eine bequeme Notation:

Definition (Hausdorffs $\sigma\delta$ -Notation)

Sei M ein Mengensystem. Dann definieren wir M_σ und M_δ durch:

$$M_\sigma = \{ \bigcup A \mid A \subseteq M, A \text{ abzählbar} \},$$

$$M_\delta = \{ \bigcap A \mid A \subseteq M, A \text{ abzählbar}, A \neq \emptyset \}.$$

[Merkhilfe: σ für Summe, δ für Durchschnitt.]

Weiter sei

$$\mathcal{G} = \{ U \subseteq \mathbb{R} \mid U \text{ offen} \}, \quad \mathcal{F} = \{ A \subseteq \mathbb{R} \mid A \text{ abgeschlossen} \},$$

[In diesen klassischen Bezeichnungen von Hausdorff steht $\mathcal{G}$ für Gebiet, $\mathcal{F}$ für ‚fermé‘.]

Wir definieren nun durch Rekursion über α mit $1 \leq \alpha < \omega_1$ die Mengen $\Sigma_\alpha \subseteq \mathcal{P}(\mathbb{R})$ und $\Pi_\alpha \subseteq \mathcal{P}(\mathbb{R})$ wie folgt:

$$\Sigma_1 = \mathcal{G}, \quad \Pi_1 = \mathcal{F},$$

$$\Sigma_\alpha = (\bigcup_{\beta < \alpha} \Pi_\beta)_\sigma, \quad \Pi_\alpha = (\bigcup_{\beta < \alpha} \Sigma_\beta)_\delta \quad \text{für } 2 \leq \alpha < \omega_1.$$

Die Notation und der Start bei 1 anstatt bei 0 ist durch die Komplexitätsnotationen der mathematischen Logik motiviert, was uns hier nicht zu kümmern braucht.

Übung

Zeigen Sie:

- (i) $\Sigma_\alpha \subseteq \Sigma_\beta, \Pi_\alpha \subseteq \Pi_\beta$ für alle $1 \leq \alpha \leq \beta < \omega_1$.
- (ii) $\Pi_{\alpha+1} = (\Sigma_\alpha)_\delta, \Sigma_{\alpha+1} = (\Pi_\alpha)_\sigma$ für alle $1 \leq \alpha < \omega_1$.
- (iii) $\Sigma_\alpha = (\Sigma_\alpha)_\sigma, \Pi_\alpha = (\Pi_\alpha)_\delta$ für alle $1 \leq \alpha < \omega_1$.
- (iv) $\Sigma_\alpha = \{ \mathbb{R} - A \mid A \in \Pi_\alpha \}, \Pi_\alpha = \{ \mathbb{R} - A \mid A \in \Sigma_\alpha \}$ für alle $1 \leq \alpha < \omega_1$.

Setzt man noch $\Delta_\alpha = \Sigma_\alpha \cap \Pi_\alpha$ für $1 \leq \alpha < \omega_1$, so ergibt sich folgendes Bild:

$$\begin{array}{cccccccc}
 \Sigma_1 & \Sigma_2 & \Sigma_3 & \Sigma_4 & \dots & \Sigma_\omega & \Sigma_{\omega+1} & \dots \\
 \Delta_1 & \Delta_2 & \Delta_3 & \Delta_4 & \dots & \Delta_\omega & \Delta_{\omega+1} & \dots \\
 \Pi_1 & \Pi_2 & \Pi_3 & \Pi_4 & \dots & \Pi_\omega & \Pi_{\omega+1} & \dots
 \end{array}$$

Weiter rechts stehende Mengen sind immer Obermengen von allen auf allen drei Ebenen weiter links stehenden Mengen. (Man kann durch ein nichttriviales Diagonalargument zeigen, daß alle Inklusionen echt sind.) Die dritte Ebene entsteht aus der ersten durch Komplementbildung in $\mathbb{R}$ und umgekehrt. Die mittlere Ebene entsteht durch Schnittbildung. Der Leser verfolge den symmetrischen Zickzackaufbau der ersten Stufen der Hierarchie mittels (ii) der Übung.

Es gilt z.B. $\Sigma_2 = \mathcal{F}_\sigma$, $\Sigma_3 = \mathcal{G}_{\delta\sigma}$, $\Sigma_4 = \mathcal{F}_{\sigma\delta\sigma}$,

Definition (*Borel-Hierarchie und Borelsche σ -Algebra*)

$\langle (\Sigma_\alpha, \Pi_\alpha) \mid 1 \leq \alpha < \omega_1 \rangle$ heißt die *Borel-Hierarchie* von $\mathbb{R}$.

$\mathfrak{B} = \mathfrak{B}(\mathbb{R}) = \bigcup_{\alpha < \omega_1} \Sigma_\alpha = \bigcup_{\alpha < \omega_1} \Pi_\alpha$ heißt die *Borelsche σ -Algebra über $\mathbb{R}$* .

$A \subseteq \mathbb{R}$ heißt eine *Borelmenge*, falls $A \in \mathfrak{B}$.

Der Ausdruck *σ -Algebra* ruht auf den guten Abschlußeigenschaften von $\mathfrak{B}$. Diese und eine Charakterisierung von $\mathfrak{B}$ gibt die folgende Übung.

Übung

(i) Es gilt $\mathfrak{B}_\sigma = \mathfrak{B}_\delta = \mathfrak{B} = \{ \mathbb{R} - A \mid A \in \mathfrak{B} \}$.

(ii) $\mathfrak{B}$ ist das $\subseteq$ -kleinste $\mathfrak{A} \supseteq \mathcal{G} \cup \mathcal{F}$ mit $\mathfrak{A}_\sigma, \mathfrak{A}_\delta \subseteq \mathfrak{A}$.

[zu (i): Ist $A \subseteq \mathbb{R}$ abzählbar, so existiert ein $\alpha < \omega_1$ mit $A \subseteq \Sigma_\alpha$.

zu (ii): Zeige durch Induktion $\Sigma_\alpha \subseteq \mathfrak{A}$, $\Pi_\alpha \subseteq \mathfrak{A}$.]

Die Aussage (ii) ist der Grund, warum die Rekursion bei ω_1 aufhört. Nach ω_1 -vielen Iterationen entsteht durch Anwendung der $\delta\sigma$ -Operationen nichts Neues mehr.

Beweis des Wohlordnungssatzes durch Abtragen

Drei Aspekte des transfiniten Abtragens einer Menge sind mittlerweile deutlich geworden:

- (1) Das Abtragen verläuft entlang von Wohlordnungen bzw. Ordinalzahlen.
- (2) Wir brauchen eine Funktion, die uns an jeder Stelle ein Element der Restmenge liefert.
- (3) Wir müssen sicherstellen, daß die Struktur, entlang derer wir abtragen, lang genug ist, damit wir mit dem Abtragen auch wirklich fertig werden. Die reellen Zahlen etwa kann man nicht in abzählbar vielen Schritten abtragen.

Im obigen ersten Beispiel einer transfiniten Rekursion haben wir implizit den Wohlordnungssatz für (3) verwendet und Ordinalzahlen für (1) zugrunde gelegt. Wir zeigen nun, daß sich dies vermeiden läßt, und erhalten so einen neuen Beweis des Wohlordnungssatzes. Zudem genügen uns Wohlordnungen als zugrundeliegende Struktur, es ist nicht unbedingt notwendig, mit Ordinalzahlen zu arbeiten.

Zunächst ist klar, daß Induktion und Rekursion entlang von Wohlordnungen $\langle M, < \rangle$ durchgeführt werden können, nicht nur entlang der Ordinalzahlen. Die Sätze hierzu lauten:

Satz (*Induktion entlang einer Wohlordnung*)

Sei $\langle M, < \rangle$ eine Wohlordnung, und sei $\mathcal{E}(x)$ eine Aussage.

Für alle $x \in M$ gelte:

(+) Gilt $\mathcal{E}(y)$ für alle $y \in M$ mit $y < x$, so gilt $\mathcal{E}(x)$.

Dann gilt $\mathcal{E}(x)$ für alle $x \in M$.

Satz (*Rekursion entlang einer Wohlordnung*)

Sei $\langle M, < \rangle$ eine Wohlordnung, und sei $\mathcal{F}$ eine Operation auf V .

Dann existiert genau eine Funktion G auf M , sodaß gilt:

(+) Für alle $x \in M$ ist $G(x) = \mathcal{F}(G \upharpoonright M_x)$.

Die Beweise sind völlig analog zu den entsprechenden Sätzen für die Ordinalzahlen.

Im folgenden Beweis des Wohlordnungssatzes spielt der Satz von Hartogs eine entscheidende Rolle: Er garantiert ohne Rückgriff auf den Wohlordnungssatz für jede Menge M die Existenz einer Wohlordnung, die lang genug ist, um zu garantieren, daß das Abtragen terminiert.

Beweis des Wohlordnungssatzes durch Abtragen

Sei M eine Menge. Sei $\mathcal{H}(M) = \langle W, < \rangle$ die Hartogswohlordnung von M .

Wir tragen M entlang $\mathcal{H}(M)$ rekursiv ab.

Sei hierzu $f : \mathcal{P}(M) - \{ \emptyset \} \rightarrow M$ eine Funktion mit $f(A) \in A$ für alle nichtleeren $A \subseteq M$. Wir definieren durch Rekursion über $x \in W$ entlang $\mathcal{H}(M)$ solange möglich $m_x \in M$ durch:

$$m_x = f(M - \{ m_y \mid y < x \}).$$

Die Definition bricht an einer Stelle $x^* \in W$ ab.

Andernfalls ist $\langle m_x \mid x \in W \rangle : W \rightarrow M$ injektiv,

im Widerspruch zu $\langle W, < \rangle$ Hartogswohlordnung von M .

Dann ist aber $\langle m_x \mid x < x^* \rangle : W_{x^*} \rightarrow M$ bijektiv,

– und induziert eine Wohlordnung auf M .

Wie in den Beweisen von Zermelo fixieren wir ganz zu Beginn eine Resthaufenfunktion f , und tragen dann M entlang der Hartogswohlordnung von M ab.

Der Satz von Hartogs (1915) war es, der Cantor gefehlt hat, um ein klares Argument für das Fertigwerden zu liefern. Der Beweis des Satzes von Hartogs selbst ist dabei trickreich, aber elementar: Der Wohlordnungssatz oder ein „ein ...“ wird nicht gebraucht. Im obigen Beweis wird ein Auswahlakt dann zur Gewinnung von f eingesetzt.

Man kann die Verwendung des Wohlordnungssatzes auch vermeiden, indem man das Fertigwerden mit einem Argument zeigt, das wir im Beweis des Repräsentationssatzes schon verwendet haben: Wir tragen M entlang der Ordinalzahlen ab wie im Beispiel (1) oben. *Annahme*, die Rekursion bricht nicht ab. Sei dann

$M' = \{x \in M \mid x = x_\alpha \text{ für eine Ordinalzahl } \alpha\}$. Weiter sei die Funktion g auf M' definiert durch $g(x) = \text{„das } \alpha \text{ mit } x = x_\alpha\text{“}$, und sei $A = \text{rng}(g)$. Dann ist A eine Menge, da A die Mächtigkeit von $M' \subseteq M$ hat. Andererseits ist jede Ordinalzahl α ein Element von A , da $x_\alpha \in M'$ definiert ist. Also ist auch $\sup(A) + 1 \in A$, *Widerspruch*.

Es ist bemerkenswert, daß Cantor dieses Argument bereits einige Zeit vor der Jahrhundertwende entwickelt hatte und damit den Wohlordnungssatz bewies (Cantor 1991, Briefe; siehe die Auszüge am Ende 2.13). Lediglich das Prinzip „kleine Vielheiten sind Mengen“ muß akzeptiert werden, um die Behauptung „ $A = \text{rng}(g)$ ist eine Menge“ zu rechtfertigen. Dieses Prinzip verwendet die Axiomatik nach Cantor einfach als Axiom (Ersetzungssaxiom). Es wird zudem ohnehin im Limeschritt des allgemeinen Rekursionssatzes verwendet – sowohl im Rekursionssatz für Ordinalzahlen als auch im Rekursionssatz entlang Wohlordnungen.

Mit jedem der beiden Argumente kann das Fertigwerden (auch innerhalb der formal axiomatischen Mengenlehre) ohne Zuhilfenahme des Wohlordnungssatzes bewiesen werden, denn die Entwicklung der Ordinalzahlen und der Beweis des Rekursionssatzes benötigen den Wohlordnungssatz nicht.

Auf der anderen Seite ist es bemerkenswert, daß der Beweis des Wohlordnungssatzes von Zermelo ohne das „Prinzip der kleinen Vielheiten“ bzw. seinem formalen Analogon, dem Ersetzungssaxiom, auskommt, und in diesem Sinne elementarer erscheint als der Beweis der Wohlordnenbarkeit einer Menge durch Abtragen entlang der Ordinalzahlen mit Hilfe eines starken Rekursionssatzes. Man kann aber auch ohne das „Prinzip der kleinen Vielheiten“ den Rekursionssatz für gewisse Operationen $\mathcal{F}$ beweisen, und für den Fall des rekursiven Abtragens mit $\mathcal{F}(G) = f(M - \text{rng}(G))$ ist dies in der Tat möglich. Nicht möglich ist dies z. B. für $\mathcal{F}(M) = \mathcal{P}(M)$ für Mengen M . Hier wird im Limeschritt gebraucht, daß die „kleine Vielheit“ $\{\mathcal{P}^n(M) \mid n \in \mathbb{N}\}$ für jedes M eine Menge ist.

Wir beenden das Kapitel mit John von Neumanns Beweis des Rekursionssatzes. Die Notation des folgenden Zitats wurde ausnahmsweise an die heutige Schreibweise angepaßt, um eine bessere Lesbarkeit zu erreichen. Von Neumann schreibt P, Q für Ordinalzahlen, $M(x; \mathcal{E}(x))$ für $\{x \mid \mathcal{E}(x)\}$, usw.

John von Neumann über den Rekursionssatz

„Von dieser Stelle an ist es leicht, die Theorie der Ordnungszahlen weiter zu entwickeln... Die ‚Definition durch Transfinite Induktion‘ ist allerdings nur dann zulässig, wenn der folgende Satz bewiesen ist:

„ $\mathcal{F}(x)$ sei eine Funktion, die für alle Mengen von Dingen eines Bereiches $B[V]$ definiert ist und deren Werte stets Dinge des Bereichs B sind. Es gibt dann eine und nur eine Funktion $\Phi(\alpha)$, die für alle Ordnungszahlen α definiert ist und deren Werte stets Dinge des Bereichs B sind, mit der Eigenschaft, daß für alle Ordnungszahlen α

$$\Phi(\alpha) = \mathcal{F}(\{\Phi(\beta) \mid \beta < \alpha\})$$

ist.“

Der Beweis dieses überhaupt nicht selbstverständlichen Satzes ist aber unschwer zu erbringen.⁴⁾ ...

[Fußnote] 4) Der Beweis des Satzes ... verläuft etwa folgendermaßen.

a) α sei eine Ordnungszahl. Gibt es dann eine Funktion $\Psi(\beta)$, die für alle Ordnungszahlen $\beta < \alpha$ definiert ist, und die Eigenschaft

$$\Psi(\beta) = \mathcal{F}(\{ \Psi(\gamma) \mid \gamma < \beta \})$$

besitzt [für alle $\beta < \alpha$]? Und wieviele solche gibt es?

b) Wenn es überhaupt eine gibt, so gibt es eine einzige. In diesem Falle heiße α ‚normal‘, und Ψ heiße Ψ_α . Ferner sei für ein normales α

$$\Phi(\alpha) = \mathcal{F}(\{ \Psi_\alpha(\beta) \mid \beta < \alpha \}).$$

c) Wenn α normal ist, so ist jedes $\beta < \alpha$ normal, und es ist für alle $\gamma < \beta$

$$\Psi_\beta(\gamma) = \Psi_\alpha(\gamma) \text{ und } \Phi(\beta) = \Psi_\alpha(\beta).$$

d) Wenn alle $\beta < \alpha$ normal sind, so ist auch α normal. (Es ist $\Psi_\alpha(\beta) = \Phi(\beta)$.)

e) Alle α sind normal. Aus d) folgt unmittelbar

$$\Phi(\alpha) = \mathcal{F}(\{ \Psi_\alpha(\beta) \mid \beta < \alpha \}) = \mathcal{F}(\{ \Phi(\beta) \mid \beta < \alpha \}),$$

also ist $\Phi(\alpha)$ die gewünschte Funktion.

f) Es gibt nur eine derartige Funktion.“

(John von Neumann 1923, „Zur Einführung der transfiniten Zahlen“, Notation modernisiert)

8. Typen linearer Ordnungen und ihre Arithmetik

Der Fundamentalsatz über Wohlordnungen hat gezeigt, daß die Wohlordnungen eine einfache und klare Struktur besitzen. Sie sind zwar reichhaltig vorhanden, aber jede Wohlordnung hat eine bestimmte Länge, die die Ordnung bis auf Isomorphie charakterisiert.

Allgemeine lineare Ordnungen sind wesentlich komplizierter. Für einige bestimmte lineare Ordnungen kann man aber interessante Charakterisierungen angeben. Insbesondere gilt dies für die Ordnungen der rationalen und der reellen Zahlen, die wir in Kapitel 10 untersuchen werden. Zuvor beschäftigen wir uns allgemein mit den Typen linearer Ordnungen und führen natürliche arithmetische Operationen für diese Ordnungstypen ein.

Zunächst erweitern wir den Begriff der Ähnlichkeit auf beliebige lineare Ordnungen.

Definition (*ähnliche lineare Ordnungen*)

Zwei lineare Ordnungen $\langle M, < \rangle$ und $\langle N, < \rangle$ heißen *ähnlich* oder *ordnungsisomorph*, falls eine Bijektion $f : M \rightarrow N$ existiert derart, daß für alle $x, y \in M$ gilt:

$$x < y \text{ gdw } f(x) < f(y).$$

Ein solches f heißt ein *Ordnungsisomorphismus* zwischen $\langle M, < \rangle$ und $\langle N, < \rangle$. Wir schreiben $\langle M, < \rangle \equiv \langle N, < \rangle$, falls $\langle M, < \rangle$ und $\langle N, < \rangle$ ähnlich sind.

Es genügt wieder die nun schon bekannte schwächere Form:

Übung

Sei $f : M \rightarrow N$ surjektiv, und für alle $x, y \in M$ gelte: $x < y$ folgt $f(x) < f(y)$. Dann ist f ein Ordnungsisomorphismus.

Cantor (1895): „Zwei geordnete Mengen M und N nennen wir ‚ähnlich‘, wenn sie sich gegenseitig eindeutig einander so zuordnen lassen, daß wenn m_1 und m_2 irgend zwei Elemente von M , n_1 und n_2 die entsprechenden Elemente von N sind, alsdann immer die Rangbeziehung von m_1 zu m_2 innerhalb M dieselbe ist wie die von n_1 und n_2 innerhalb N ...“

In der Regel kann man keine Eindeutigkeit des Ordnungsisomorphismus zwischen ähnlichen linearen Ordnungen erwarten. Für $\langle \mathbb{Z}, < \rangle$ ist zum Beispiel für

jedes $a \in \mathbb{Z}$ die Translation $f_a : \mathbb{Z} \rightarrow \mathbb{Z}$ mit $f_a(z) = z + a$ für $z \in \mathbb{Z}$ ein Ordnungs-
isomorphismus von $\mathbb{Z}$ in sich selbst.

Wir denken uns lineare Ordnungen wieder mit einem Ordnungstyp versehen
derart, daß zwei lineare Ordnungen genau dann den gleichen Ordnungstyp ha-
ben, wenn sie ähnlich sind:

$\text{o.t.}(\langle M, < \rangle) = \text{o.t.}(\langle N, < \rangle)$ gdw $\langle M, < \rangle$ und $\langle N, < \rangle$ sind ordnungsisomorph.

Formal kann $\text{o.t.}(\langle M, < \rangle)$ mit Hilfe einer Trunkierung aller zu $\langle M, < \rangle$ ordnungsisomor-
phen Ordnungen nach dem Rang definiert werden (vgl. Kapitel 7, Diskussion der V_α -
Hierarchie).

Wir verwenden vorwiegend kleine griechische Buchstaben für Ordnungstypen. Für wichtige Ordnungstypen gibt es feste Bezeichnungen. Eine kleine
Tabelle von Ordnungstypen ist:

Definition (die Ordnungstypen n für $n \in \mathbb{N}$, ω , ζ , η und θ)

Wir setzen:

$$n = \text{o.t.}(\langle \{0, 1, \dots, n-1\}, < \rangle) \quad \text{für } n \in \mathbb{N},$$

$$\omega = \text{o.t.}(\langle \mathbb{N}, < \rangle),$$

$$\zeta = \text{o.t.}(\langle \mathbb{Z}, < \rangle),$$

$$\eta = \text{o.t.}(\langle \mathbb{Q}, < \rangle),$$

$$\theta = \text{o.t.}(\langle \mathbb{R}, < \rangle),$$

wobei $<$ jeweils die übliche Ordnung bezeichnet.

Ein Typus α heißt *abzählbarer Ordnungstyp*, falls $\alpha = \text{o.t.}(\langle M, < \rangle)$ für eine lineare
Ordnung $\langle M, < \rangle$ mit abzählbarem Träger M . Eine einfache Aussage über die
Anzahl der abzählbaren Typen ist:

Übung

Es gibt höchstens $\mathbb{R}$ -viele abzählbare Ordnungstypen,
d.h. es gibt eine Menge $\mathbb{O}$ von linearen Ordnungen mit:

$$(i) \quad |\mathbb{O}| \leq |\mathbb{R}|.$$

- (ii) Für jede abzählbare lineare Ordnung $\langle M, < \rangle$ existiert ein
 $\langle N, < \rangle \in \mathbb{O}$ mit $\langle M, < \rangle \equiv \langle N, < \rangle$.

[Ist $\langle M, < \rangle$ abzählbar, so ist $<$ bis auf Isomorphie eine Teilmenge von $\mathbb{N}^2$.]

Wir werden später sehen, daß es genau $\mathbb{R}$ -viele abzählbare Ordnungstypen gibt, eine
Beobachtung von Cantor (mindestens $\mathbb{R}$ -viele) sowie Bernstein und Hausdorff (für höch-
stens $\mathbb{R}$ -viele). Vgl. hierzu auch [Hausdorff 2002, Anm. 37].

Zu einer gegebenen linearen Ordnung kann man eine neue Ordnung erhalten,
indem man die Ordnung durch Vertauschen von „größer“ und „kleiner“ um-
kehrt:

Definition (*inverse Ordnungen und α^**)

Sei $\langle M, < \rangle$ eine lineare Ordnung.

Dann heißt $\langle M, <^* \rangle$ die zu $\langle M, < \rangle$ *inverse Ordnung*, wobei gilt:

$$a <^* b \text{ falls } b < a \quad \text{für alle } a, b \in M.$$

Ist α ein Ordnungstyp, so ist α^* definiert durch

$$\alpha^* = \text{o. t.}(\langle M, <^* \rangle),$$

wobei $\langle M, < \rangle$ eine lineare Ordnung ist mit $\alpha = \text{o. t.}(\langle M, < \rangle)$.

Übung

- (i) $\alpha^{**} = \alpha$ für alle Ordnungstypen α ,
- (ii) $n^* = n$ für alle $n \in \mathbb{N}$,
- (iii) $\omega^* \neq \omega$, $\zeta^* = \zeta$, $\eta^* = \eta$, $\theta^* = \theta$.

Cantor (1895): „Werden in einer geordneten Menge M alle Rangbeziehungen ihrer Elemente umgekehrt, so daß überall aus dem ‚niedriger‘ ein ‚höher‘ und aus dem ‚höher‘ ein ‚niedriger‘ wird, so erhält man wieder eine geordnete Menge, die wir mit

$*M$

bezeichnen und die ‚inverse‘ von M nennen wollen.

Den Ordnungstypus von *M bezeichnen wir, wenn $\alpha = \bar{M} [= \text{o. t.}(\langle M, < \rangle)]$ ist, mit

$${}^*\alpha.$$

Es kann vorkommen, daß ${}^*\alpha = \alpha$, wie zum Beispiel bei den endlichen Typen oder beim Typus der Menge $\mathbb{R}$ aller rationalen Zahlen ...“

Ordnungstypen lassen sich in natürlicher Weise addieren und multiplizieren. Dabei treten im Vergleich zur endlichen Arithmetik einige Besonderheiten auf.

Die Addition von Ordnungstypen

Sind $\langle M, < \rangle$ und $\langle N, < \rangle$ zwei Ordnungen, so können wir ohne Einschränkung stets annehmen, daß die Ordnungen disjunkt sind – d. h. daß $M \cap N = \emptyset$ gilt –, solange es uns nur auf den Typ ankommt. Denn für eine lineare Ordnung und ein beliebiges Objekt i sei $M^i = M \times \{i\}$. Wir definieren $\langle M^i, < \rangle$ durch

$$(x, i) < (y, i) \text{ falls } x < y \quad \text{für alle } (x, i), (y, i) \in M^i.$$

Dann ist $\langle M^i, < \rangle$ eine lineare Ordnung und es gilt $\text{o. t.}(\langle M, < \rangle) = \text{o. t.}(\langle M^i, < \rangle)$. Für zwei Ordnungen $\langle M, < \rangle$ und $\langle N, < \rangle$ sind dann $\langle M^0, < \rangle$ und $\langle N^1, < \rangle$ voneinander jeweils gleichen Ordnungstyp und es gilt $M^0 \cap N^1 = \emptyset$.

Nach dieser Vorbemerkung können wir nun die Addition zweier linearer Ordnungen definieren.

Definition (*Addition zweier linearer Ordnungen*)

Seien $\langle M, < \rangle$ und $\langle N, < \rangle$ lineare Ordnungen mit $M \cap N = \emptyset$.

Dann ist die *Summe* $\langle S, < \rangle$ von $\langle M, < \rangle$ und $\langle N, < \rangle$, in Zeichen

$$\langle S, < \rangle = \langle M, < \rangle + \langle N, < \rangle$$

definiert durch:

$$S = M \cup N,$$

$$a < b \quad \text{falls} \quad \begin{array}{l} a, b \in M \text{ und } a < b \quad \text{oder} \\ a, b \in N \text{ und } a < b \quad \text{oder} \\ a \in M, b \in N \end{array} \quad \text{für } a, b \in S.$$

M	$+$	N
-----	-----	-----

Die Ordnung $\langle N, < \rangle$ wird also einfach an die Ordnung $\langle M, < \rangle$ rechts angehängt – mit der Konvention, daß ein größeres Element einer Ordnung weiter rechts steht als ein kleineres. In der Ordnung $\langle M, < \rangle + \langle N, < \rangle$ ist jedes Element von N größer als jedes Element von M ; sind zwei Elemente x, y aus $M \cup N$ beide im linken oder beide im rechten Teil, so werden sie nach den bereits vorhandenen Ordnungen von M und N miteinander verglichen.

Cantor (1895): „Die Vereinigungsmenge (M, N) [hier: disjunkte Vereinigung] zweier Mengen M und N läßt sich, wenn die letzteren geordnet sind, selbst als eine geordnete Menge auffassen, in welcher die Rangbeziehung der Elemente von M unter einander, ebenso die Rangbeziehung der Elemente von N unter einander dieselben wie in M resp. N geblieben sind, dagegen alle Elemente von M niedrigeren Rang als alle Elemente von N haben.“

Die Addition für Ordnungstypen wird mit Hilfe der Addition von linearen Ordnungen wie folgt definiert:

Definition (*Addition zweier Ordnungstypen*)

Seien α, β Ordnungstypen.

Dann ist die *Summe* von α und β , in Zeichen $\alpha + \beta$, definiert durch:

$$\alpha + \beta = \text{o. t.}(\langle M, < \rangle + \langle N, < \rangle),$$

wobei $\langle M, < \rangle$ und $\langle N, < \rangle$ zwei disjunkte Ordnungen sind mit

$$\alpha = \text{o. t.}(\langle M, < \rangle) \text{ und } \beta = \text{o. t.}(\langle N, < \rangle).$$

Man sieht leicht, daß diese Definition nicht von der Wahl von $\langle M, < \rangle$ und $\langle N, < \rangle$ abhängt. Für Ordinalzahlen α stimmt die neue Definition von $\alpha + 1$ mit der alten überein.

Übung

Für alle Ordnungstypen α, β, γ gilt $(\alpha + \beta) + \gamma = \alpha + (\beta + \gamma)$.

Wir können also Ausdrücke wie $\alpha + \beta + \gamma$, $\alpha + \beta + \gamma + \delta$, usw. unzweideutig ohne Klammern schreiben.

Übung

- (i) $0 + \alpha = \alpha + 0$ für alle Ordnungstypen α ,
- (ii) $\omega + 1 \neq \omega$, $1 + \omega^* \neq \omega^*$;
 $n + \omega = \omega$, $\omega^* + n = \omega^*$ für alle $n \in \mathbb{N}$,
- (iii) $(\alpha + \beta)^* = \beta^* + \alpha^*$ für alle Ordnungstypen α, β ,
- (iv) $\omega^* + n + \omega = \zeta$ für alle $n \in \mathbb{N}$.

Die Fälle $\alpha + \beta = \beta$ und $\beta + \alpha = \beta$ sind also für $\alpha \neq 0$ möglich und Kommutativität, also $\alpha + \beta = \beta + \alpha$, ist im allgemeinen nicht richtig.

Übung

Seien α, β Ordinalzahlen, $\alpha \leq \beta$.

Dann existiert genau eine Ordinalzahl γ mit $\alpha + \gamma = \beta$.

Die Multiplikation von Ordnungstypen

Die Multiplikation von natürlichen Zahlen kann man sich durch Rechtecke veranschaulichen. Auch für beliebige lineare Ordnungen kann nun eine natürliche Multiplikation mit Hilfe des kartesischen Produkts definiert werden.

Definition (Multiplikation zweier linearer Ordnungen)

Seien $\langle M, < \rangle$ und $\langle N, < \rangle$ lineare Ordnungen.

Dann ist das *Produkt* $\langle P, < \rangle$ von $\langle M, < \rangle$ und $\langle N, < \rangle$, in Zeichen

$$\langle P, < \rangle = \langle M, < \rangle \cdot \langle N, < \rangle$$

definiert durch:

$$P = M \times N,$$

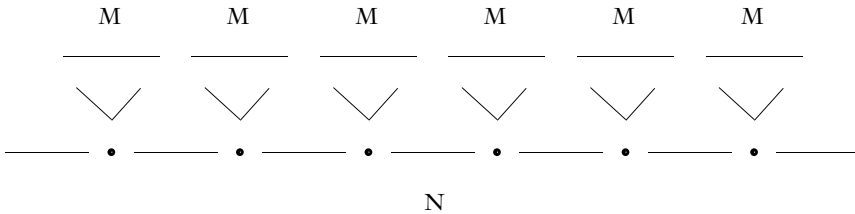
$$(a_1, b_1) < (a_2, b_2) \text{ falls } b_1 < b_2 \text{ oder } b_1 = b_2 \text{ und } a_1 < a_2 \quad \text{für } (a_1, b_1), (a_2, b_2) \in P.$$

N

M

In der Skizze sei M von links nach rechts und N von unten nach oben geordnet. Dann kann man die Ordnung auf $M \times N$ wie folgt beschreiben: Punkte auf höher liegenden Zeilen sind immer größer als Punkte auf niedrigeren Zeilen. Innerhalb jeder Zeile werden die Punkte gemäß M geordnet.

Eine andere Möglichkeit, sich die Produktordnung auf $M \times N$ vorzustellen, ist die folgende: Für jedes Element von N setzen wir eine Kopie der Ordnung $\langle M, < \rangle$ ein. Elemente werden also zu Ordnungen; die neue Menge wird mit der Priorität der Einsetzstelle und innerhalb einer Einsetzung gemäß M geordnet:



Anders: Wir betrachten alle Punkte von N durch ein Mikroskop, und lösen dadurch jeden Punkt in eine zu M ähnliche Ordnung auf.

Cantor (1895): „Aus zwei geordneten Mengen M und N mit den Typen α und β läßt sich eine geordnete Menge S dadurch herstellen, daß in N an die Stelle jedes Elementes n eine geordnete Menge M_n substituiert wird, welche denselben Typus α wie M hat ..., und daß über die Rangordnung in S ... folgende Bestimmungen getroffen werden:

- 1) je zwei Elemente von S , welche einer und derselben Menge M_n angehören, behalten in S dieselbe Rangbeziehung wie in M_n .
 - 2) je zwei Elemente von S , welche zwei verschiedenen Mengen M_{n_1} und M_{n_2} angehören, erhalten in S die Rangbeziehung, welche n_1 und n_2 in N haben.“
-

Eine wichtige Beobachtung ist, daß Produkte aus der Klasse der Wohlordnungen nicht herausführen:

Übung

Das Produkt zweier Wohlordnungen ist eine Wohlordnung.

Die Multiplikation für Ordnungstypen ist nun analog zur Addition definiert:

Definition (Multiplikation zweier Ordnungstypen)

Seien α, β Ordnungstypen.

Dann ist das *Produkt* von α und β , in Zeichen $\alpha \cdot \beta$, definiert durch:

$$\alpha \cdot \beta = \text{o.t.}(\langle M, < \rangle \cdot \langle N, < \rangle),$$

wobei $\langle M, < \rangle$ und $\langle N, < \rangle$ zwei Ordnungen sind mit

$$\alpha = \text{o.t.}(\langle M, < \rangle) \text{ und } \beta = \text{o.t.}(\langle N, < \rangle).$$

Wieder hängt diese Definition nicht von der Wahl von $\langle M, < \rangle$ und $\langle N, < \rangle$ ab.

Übung

- (i) Seien α, β von 0 verschiedene Ordnungstypen. Dann ist $\alpha \cdot \beta$ genau dann eine Ordinalzahl, falls α und β Ordinalzahlen sind.
- (ii) Es gilt $\alpha \cdot (\beta + \gamma) = \alpha \cdot \beta + \alpha \cdot \gamma$ für alle Ordnungstypen α, β, γ .
- (iii) Im allgemeinen ist $(\alpha + \beta) \cdot \gamma = \alpha \cdot \gamma + \beta \cdot \gamma$ falsch.
- (iv) Es gilt $\alpha \cdot (\beta \cdot \gamma) = (\alpha \cdot \beta) \cdot \gamma$ für alle Ordnungstypen α, β, γ .

Insbesondere haben wir: $\omega \cdot 2 = \omega + \omega$, $2 \cdot \omega = \omega$. Die Multiplikation ist also wie die Addition nicht mehr kommutativ. Wegen der Assoziativität können wir wieder Klammern in Produkten weglassen. Wir schreiben auch $\alpha \cdot \beta$ für $\alpha \cdot \beta$.

Allgemeine Summen von Ordnungstypen

Definition (*Summe von linearen Ordnungen über eine lineare Ordnung*)

Sei $\langle N, < \rangle$ eine lineare Ordnung, und seien $\langle M_x, < \rangle$ lineare Ordnungen für $x \in N$. Sei $S = \bigcup_{x \in N} M_x \times \{x\}$. Für $(a, x), (b, y) \in S$ setzen wir:

$(a, x) < (b, y)$ falls $x < y$ in $\langle N, < \rangle$ oder
 $x = y$ und $a < b$ in $\langle M_x, < \rangle$.

$\langle S, < \rangle$ heißt die *Summe der Ordnungen* $\langle M_x, < \rangle$ über $\langle N, < \rangle$, in Zeichen

$$\langle S, < \rangle = \sum_{x \in N}^{\langle N, < \rangle} \langle M_x, < \rangle.$$

Eine derartige Summe ist offenbar eine lineare Ordnung.

Wir lassen den oberen Index $\langle N, < \rangle$ an der Summe oft weg, wenn die Ordnung auf N klar ist, und schreiben dann übersichtlicher $\sum_{x \in N} \langle M_x, < \rangle$.

Der Typ der Summe hängt nur von den Typen der Ordnungen $\langle M_x, < \rangle$ ab, und wir definieren:

Definition (*Summe von Ordnungstypen*)

Sei $\langle N, < \rangle$ eine lineare Ordnung, und seien α_x Ordnungstypen für $x \in N$. Dann ist die *Summe* von $\langle \alpha_x \mid x \in N \rangle$ über $\langle N, < \rangle$, in Zeichen $\sum_{x \in N}^{\langle N, < \rangle} \alpha_x$, definiert durch:

$$\sum_{x \in N}^{\langle N, < \rangle} \alpha_x = \text{o.t.}(\sum_{x \in N}^{\langle N, < \rangle} \langle M_x, < \rangle),$$

wobei $\langle M_x, < \rangle$ für $x \in N$ beliebige lineare Ordnungen sind mit
 $\alpha_x = \text{o.t.}(\langle M_x, < \rangle)$ für $x \in N$.

Ein Unterschied zu den Definitionen von $\alpha + \beta$ und $\alpha \cdot \beta$ ist hier, daß in der Summe links noch eine lineare Ordnung $\langle N, < \rangle$ stehen bleibt, und wir nicht nur mit dem Typ $\beta = \text{o.t.}(\langle N, < \rangle)$ auskommen; die Typen α_x brauchen einen Index.

Es gilt z.B. $\sum_{n \in \mathbb{N}} 1 = \sum_{n \in \mathbb{N}} n = \omega$. Der Leser überlegt sich leicht die folgenden Gleichungen für alle linearen Ordnungen $\langle N, < \rangle$:

- (a) $\sum_{x \in N} 1 = \text{o.t.}(\langle N, < \rangle)$,
- (b) $\sum_{x \in N} \alpha = \alpha \cdot \text{o.t.}(\langle N, < \rangle)$ für alle Ordnungstypen α .

Wir verzichten hier auf die weitere Diskussion von Rechenregeln, nutzen aber die allgemeinen Summen noch für die Konstruktion von $\mathbb{R}$ -vielen abzählbaren Ordnungstypen:

Übung

Es gibt genau $\mathbb{R}$ -viele abzählbare Ordnungstypen.

[„Höchstens“ wissen wir nach einer früheren Übung.

Für jedes $f: \mathbb{N} \rightarrow \mathbb{N}$ sei $\alpha_f = \sum_{n \in \mathbb{N}} (f(n) + \zeta)$, wobei $\zeta = \text{o.t.}(\langle \mathbb{Z}, < \rangle)$.

Dann gilt $\alpha_f \neq \alpha_g$ für alle $f, g: \mathbb{N} \rightarrow \mathbb{N}$ mit $f \neq g$.

Zeige hierzu induktiv: $\alpha_f = \alpha_g$ folgt $f(n) = g(n)$ für alle $n \in \mathbb{N}$.

Alternativ betrachte $\beta_f = \sum_{n \in \mathbb{N}} (1 + f(n) + 1 + \eta)$, wobei $\eta = \text{o.t.}(\langle \mathbb{Q}, < \rangle)$.

Die beiden zusätzlichen Einsen sind hier wichtig; im übernächsten Kapitel zeigen wir $\eta = \eta + \eta = \eta + 1 + \eta$. Offenbar aber $\eta \neq \eta + 2 + \eta$.]

Es gibt also genau $\mathbb{R}$ -viele (also $2^{\aleph_0}$ -viele) abzählbare Ordnungstypen, und es gibt genau $\aleph_1$ -viele Ordnungstypen von abzählbaren Wohlordnungen. Hausdorff und andere hatten gehofft, daß sich mit diesen Beobachtungen vielleicht das Kontinuumsproblem lösen ließe: Konstruiere eine Bijektion zwischen der Menge $T(\aleph_0) = \{ \alpha \mid \alpha \text{ ist ein abzählbarer Ordnungstyp} \}$ und ihrer Teilmenge $W(\omega_1) = \{ \alpha \mid \alpha \text{ ist Typ einer abzählbaren Wohlordnung} \}$. Hieraus würde (CH) folgen. Eine hübsche Idee, die leider nicht funktioniert.

Die Exponentiation von Ordnungstypen

Sind $\langle M, < \rangle$ und $\langle N, < \rangle$ lineare Ordnungen, so scheint die Menge

$${}^N M = \{ f \mid f: N \rightarrow M \}$$

ein natürlicher Kandidat für den Träger E einer Exponentiation

$$\langle E, < \rangle = \langle M, < \rangle^{\langle N, < \rangle}$$

zu sein – so wie $M \cup N$ Träger der Summe und $M \times N$ Träger des Produktes zweier Ordnungen waren.

Die Definition einer Ordnung auf ${}^N M$ macht aber Schwierigkeiten. Wir betrachten hierzu $M = \{0, 1\}$ und $N = \mathbb{Z}$ mit den üblichen Ordnungen. Wir definieren $f, g \in {}^{\mathbb{Z}}\{0, 1\}$ durch:

$$f(z) = \begin{cases} 0, & \text{falls } |z| \text{ gerade,} \\ 1, & \text{sonst.} \end{cases} \quad g(z) = \begin{cases} 0, & \text{falls } |z| \text{ ungerade,} \\ 1, & \text{sonst.} \end{cases}$$

Nach welchem Kriterium soll man festlegen, ob $f < g$ oder $g < f$ gelten soll? f und g sind vollkommen symmetrisch gebaut.

Ist der Exponent $\langle N, < \rangle$ aber eine Wohlordnung, so läßt sich eine natürliche lineare Ordnung für beliebige $\langle M, < \rangle$ leicht angeben.

Definition (die lexikographische Ordnung)

Seien $\langle M, < \rangle$ eine lineare Ordnung und $\langle N, < \rangle$ eine Wohlordnung. Dann ist die *lexikographische Exponentiation* von $\langle M, < \rangle$ und $\langle N, < \rangle$, in Zeichen

$$\langle E, <_{\text{lex}} \rangle = \exp_{\text{lex}}(\langle M, < \rangle, \langle N, < \rangle),$$

wie folgt definiert:

Sei $E = {}^N M$.

Für $f, g \in E$ mit $f \neq g$ sei $\delta(f, g) = \text{„das kleinste } x \in N \text{ mit } f(x) \neq g(x)\text{“}$.

Wir setzen dann für $f, g \in E$:

$f <_{\text{lex}} g$ falls $f \neq g$ und $f(\delta) < g(\delta)$ für $\delta = \delta(f, g)$.

$<_{\text{lex}}$ heißt die *lexikographische Ordnung* auf ${}^N M$ (bzgl. $\langle M, < \rangle$ und $\langle N, < \rangle$).

Wir suchen also nach der $<_N$ -kleinsten Stelle eines Unterschieds zwischen f und g . Für $f \neq g$ existiert diese Stelle, denn die Ordnung auf N ist eine Wohlordnung. Danach vergleichen wir die Werte von f und g an dieser Stelle gemäß der Ordnung von M . Nicht völlig trivial ist die Transitivität:

Übung

$\langle E, <_{\text{lex}} \rangle$ ist eine lineare Ordnung.

Definition (die lexikographische Exponentiation zweier Ordnungstypen)

Seien α ein Ordnungstyp und β eine Ordinalzahl.

Dann ist die *lexikographische Exponentiation* von α und β , in Zeichen $\exp_{\text{lex}}(\alpha, \beta)$, definiert durch:

$$\exp_{\text{lex}}(\alpha, \beta) = \text{o.t.}(\exp_{\text{lex}}(\langle M, < \rangle, \langle N, < \rangle)),$$

wobei $\langle M, < \rangle$ und $\langle N, < \rangle$ zwei Ordnungen sind mit $\alpha = \text{o.t.}(\langle M, < \rangle)$ und $\beta = \text{o.t.}(\langle N, < \rangle)$.

Die Definition ist unabhängig von der Wahl von $\langle M, < \rangle$ und $\langle N, < \rangle$.

Mit der lexikographischen Exponentiation von $\langle M, < \rangle$ und $\langle N, < \rangle$ erhalten wir also in vielen Fällen eine lineare Ordnung auf der vollen Belegungsmenge ${}^N M$. Insbesondere ist die lexikographische Exponentiation zweier Wohlordnungen eine lineare Ordnung. Somit können wir zum Beispiel ${}^{\mathbb{N}}\{0, 1\}$ und damit $\mathbb{R}$ linear ordnen – was keinerlei Neuerung bringt. Eine Wohlordnung von $\mathbb{R}$ auf eine so einfache Weise zu erhalten wäre aber in der Tat eine Neuerung. Leider führt die lexikographische Exponentiation aber aus der Klasse der Wohlordnungen heraus, und ist daher als arithmetische Operation für die Ordinalzahlen ungeeignet.

Übung

$\exp_{\text{lex}}(2, \omega)$ ist keine Ordinalzahl.

Wir wollen nun eine andere Exponentiation für Ordinalzahlen definieren, die wieder Ordinalzahlen ergibt. Hierzu orientieren wir uns an der Rekursionsgleichung $m^{n+1} = m^n \cdot m$ der üblichen Exponentiation für natürliche Zahlen.

Definition (*die Exponentiation von Ordinalzahlen*)

Sei α eine Ordinalzahl. Wir definieren die (*gewöhnliche*) *Exponentiation* α^β für Ordinalzahlen β durch transfinite Rekursion über β wie folgt:

$$\begin{aligned}\alpha^0 &= 1, \\ \alpha^{\beta+1} &= \alpha^\beta \cdot \alpha, && \text{für } \beta \text{ Ordinalzahl,} \\ \alpha^\lambda &= \sup_{\beta < \lambda} \alpha^\beta && \text{für } \lambda \text{ Limesordinalzahl.}\end{aligned}$$

Dann ist α^β eine Ordinalzahl für alle Ordinalzahlen α und β .

In dieser rekursiven Form hat Cantor die Exponentiation von Ordinalzahlen definiert (für abzählbare Ordinalzahlen; Cantor, 1897, § 18). Zermelo schreibt in einer Anmerkung hierzu:

Zermelo (1932): „Die Einführung eines neuen *Potenzbegriffes*, der von dem früher für die Kardinalzahlen gegebenen wesentlich verschieden ist, ermöglicht erst eine formalarithmetische Theorie der transfiniten Ordnungszahlen, und zwar nicht nur in der zweiten Zahlenklasse [abzählbar unendliche Ordinalzahlen]. Die zu ihrer Einführung hier verwendete Methode der ‚Definition durch transfinite Induktion‘ (Hausdorff [1914], Kap. V, § 3) ist seitdem vorbildlich geworden für alle solchen transfiniten Konstruktionen.“

Es gelten zwei der drei üblichen Rechenregeln:

Übung

- (i) Es gilt $\alpha^{\beta+\gamma} = \alpha^\beta \cdot \alpha^\gamma$ für alle Ordinalzahlen α, β, γ .
- (ii) Es gilt $\alpha^{\beta \cdot \gamma} = (\alpha^\beta)^\gamma$ für alle Ordinalzahlen α, β, γ .
- (iii) $\alpha^\gamma \cdot \beta^\gamma = (\alpha \cdot \beta)^\gamma$ ist im allgemeinen falsch (schon für $\gamma = 2$).
[$\alpha^2 \beta^2 = \alpha \alpha \beta \beta$, $(\alpha \beta)^2 = \alpha \beta \alpha \beta$. Wähle z. B. $\alpha = \omega$, $\beta = 2$.]
- (iv) Seien β, γ Ordinalzahlen, und sei $\beta < \gamma$.
Dann gilt $\alpha^\beta < \alpha^\gamma$ für alle Ordinalzahlen $\alpha > 1$.

Die ersten Ordinalzahlen

Mit Hilfe der arithmetischen Operationen können wir nun die Struktur der ersten Ordinalzahlen leicht beschreiben. Wir betrachten zunächst die Nachfolgerbildung „+ 1“, die Addition „+ ω “ und die Multiplikation „ $\cdot \omega$ “. Iterieren wir, beginnend bei 0 bzw. 1, diese Operationen ω -oft, so erhalten wir:

$$\begin{aligned}0, 1, 2, 3, \dots, n, \dots, \omega, \\ 0, \omega, \omega 2, \omega 3, \dots, \omega n, \dots, \omega^2, \\ 1, \omega, \omega^2, \omega^3, \dots, \omega^n, \dots, \omega^\omega,\end{aligned}$$

wobei die letzte Ordinalzahl hier jeweils das Supremum der gebildeten Reihe ist.

Übung

Alle Ordinalzahlen $\alpha < \omega^\omega$ kann man eindeutig in der Form schreiben:

$$\alpha = \omega^n \beta_n + \omega^{n-1} \beta_{n-1} + \dots + \omega^1 \beta_1 + \beta_0,$$

mit $n < \omega$ und $\beta_0, \dots, \beta_n < \omega$.

Wenden wir nun in analoger Weise, beginnend bei ω , iteriert die Ordinalzahl-exponentiation „hoch ω “ an, so erhalten wir:

$$\omega, \omega^\omega, (\omega^\omega)^\omega = \omega^{\omega \cdot \omega} = \omega^{\omega^2}, (\omega^{\omega^2})^\omega = \omega^{\omega^3}, \dots, \omega^{\omega^n}, \dots, \omega^{\omega^\omega}.$$

Die Konvention ist hier und im folgenden: $\alpha^{\beta^\gamma} = \alpha^{(\beta^\gamma)}$.

Wiederholung des Verfahrens, beginnend bei ω^{ω^ω} , liefert:

$$(\omega^{\omega^\omega})^\omega = \omega^{\omega^{\omega^\omega}} = \omega^{\omega^{\omega+1}}, (\omega^{\omega^{\omega+1}})^\omega = \omega^{\omega^{\omega+2}}, \dots, \omega^{\omega^{\omega+n}}, \dots, \omega^{\omega^{\omega+\omega}} = \omega^{\omega^{\omega^2}}.$$

Wiederholung dieses Prozesses liefert:

$$\omega^{\omega^\omega}, \omega^{\omega^{\omega^2}}, \dots, \omega^{\omega^{\omega^n}}, \dots, \omega^{\omega^{\omega^\omega}} = \omega^{\omega^{\omega^2}}.$$

Wiederholung wiederum dieses Prozesses liefert:

$$\omega^{\omega^\omega}, \omega^{\omega^{\omega^2}}, \dots, \omega^{\omega^{\omega^n}}, \dots, \omega^{\omega^{\omega^\omega}}.$$

Schließlich können wir diesen Prozeß der Bildung von „ ω -Türmen“ iterieren und erhalten:

$$\omega, \omega^\omega, \omega^{\omega^\omega}, \omega^{\omega^{\omega^\omega}}, \dots, \text{„}\omega\text{-Turm der Höhe } n\text{“, } \dots,$$

wobei die Exponenten der Türme von oben nach unten geklammert werden. Für das Supremum dieser Folge hat Cantor die Bezeichnung ε_0 eingeführt.

Definition (ε_0)

Wir setzen:

$$\varepsilon_0 = \sup_{n < \omega} \alpha_n,$$

wobei α_n für $n \in \omega$ rekursiv definiert ist durch:

$$\alpha_0 = \omega,$$

$$\alpha_{n+1} = \omega^{\alpha_n} \quad \text{für } n \in \omega.$$

Als Supremum abzählbar vieler abzählbarer Ordinalzahlen ist ε_0 abzählbar. Von der ersten überabzählbaren Ordinalzahl ω_1 ist der Riese ε_0 also „fast“ ebenso weit entfernt wie die Zahl Null ...

Eine bemerkenswerte Eigenschaft von ε_0 verdient eine allgemeine Betrachtung.

Fixpunkte

Wir untersuchen Operationen $\mathcal{F}$ von den Ordinalzahlen in sich selbst, etwa $\mathcal{F}(\alpha) = \alpha + \omega$ für α Ordinalzahl.

Definition (Fixpunkte)

Sei $\mathcal{F}$ eine Operation auf den Ordinalzahlen.

Eine Ordinalzahl α heißt *Fixpunkt von $\mathcal{F}$* , falls $\mathcal{F}(\alpha) = \alpha$.

Betrachten wir die schnell wachsende Operation der Exponentiation zur Basis ω , also die Operation $\mathcal{F}$ mit $\mathcal{F}(\alpha) = \omega^\alpha$ für alle Ordinalzahlen α , so ist nicht klar, ob diese Operation überhaupt Fixpunkte besitzt. Tatsächlich sind Fixpunkte aber in großer Zahl vorhanden, und ε_0 ist der kleinste unter ihnen:

Satz (Charakterisierung von ε_0)

ε_0 ist der kleinste Fixpunkt der Ordinalzahlexponentiation zur Basis ω , d. h.:

- (i) Es gilt $\varepsilon_0 = \omega^{\varepsilon_0}$.
- (ii) Für alle $\beta < \varepsilon_0$ ist $\beta < \omega^\beta$.

Beweis

Sei α_n für $n \in \mathbb{N}$ wie oben definiert, also $\alpha_0 = \omega$, $\alpha_{n+1} = \omega^{\alpha_n}$ für $n \in \omega$.

Es gilt $\alpha_0 < \alpha_1$, und aus $\alpha_n < \alpha_{n+1}$ folgt $\alpha_{n+1} = \omega^{\alpha_n} < \omega^{\alpha_{n+1}} = \alpha_{n+2}$ für alle $n \in \omega$. Also $\alpha_n < \alpha_{n+1}$ für alle $n \in \omega$.

zu (i): Es gilt:

$$\omega^{\varepsilon_0} = \sup_{\beta < \varepsilon_0} \omega^\beta = \sup_{n \in \mathbb{N}} \omega^{\alpha_n} = \sup_{n \in \mathbb{N}} \alpha_{n+1} = \varepsilon_0.$$

Die erste Gleichung gilt nach Definition von ω^λ für Limeszahlen λ , die zweite gilt wegen $\sup \{ \beta \mid \beta < \varepsilon_0 \} = \sup \{ \alpha_n \mid n \in \mathbb{N} \}$ (und der Monotonie der Exponentiation).

zu (ii): Sei $\beta < \varepsilon_0$. Die Aussage ist klar für $\beta < \omega^1 = \omega$. Andernfalls existiert ein n mit $\alpha_n \leq \beta < \alpha_{n+1}$. Dann ist $\beta < \alpha_{n+1} = \omega^{\alpha_n} \leq \omega^\beta$.

Allgemein kann man Fixpunkte durch einen „Einholprozeß“ konstruieren:

Übung

Sei α_0 eine Ordinalzahl. Definiere rekursiv $\alpha_{n+1} = \omega^{\alpha_n}$ für $n \in \omega$.

Sei $\alpha^* = \sup_{n < \omega} \alpha_n$. Dann ist $\omega^{\alpha^*} = \alpha^*$ und für alle β mit $\alpha_0 < \beta < \alpha^*$ ist $\beta < \omega^\beta$.

Startet man hier mit $\alpha_0 = \varepsilon_0 + 1$, so erhält man $\alpha^* = \varepsilon_1$, den zweiten Fixpunkt der Exponentiation zur Basis ω , usw. Weiter ist ein Limes von Fixpunkten wieder ein Fixpunkt.

Man sieht leicht, daß die beiden folgenden Eigenschaften einer Operation $\mathcal{F}$ auf den Ordinalzahlen für derartige Fixpunkt Konstruktionen ausreichen:

- (i) $\alpha \leq \mathcal{F}(\alpha)$ für alle α , (Expansivität)
 (ii) $\mathcal{F}(\lambda) = \sup_{\alpha < \lambda} \mathcal{F}(\alpha)$ für alle Limesordinalzahlen λ . (Stetigkeit)

Diese Bedingungen sind im obigen Fall für $\mathcal{F}(\alpha) = \omega^\alpha$ erfüllt. Die erste Eigenschaft gilt insbesondere dann, wenn $\mathcal{F}$ monoton ist, d. h. wenn $\mathcal{F}(\alpha) < \mathcal{F}(\beta)$ gilt für alle $\alpha < \beta$ (vgl. den Satz von Zermelo in Kapitel 4).

Auch die Alephreihe $\mathcal{F}(\alpha) = \aleph_\alpha$ erfüllt die beiden Eigenschaften. Wir erhalten also einen „Kardinalzahl-Giganten“ $\aleph_\alpha$ mit der bemerkenswerten Eigenschaft $\alpha = \aleph_\alpha$. Anders ausgedrückt: Es gibt eine Kardinalzahl α , welche die α -te unendliche Kardinalzahl ist!

Der Additions- und Multiplikationssatz

Wir geben nun noch Beweise für $|M + M| = |M|$ und $|M \times M| = |M|$ für unendliche Mengen M mit Hilfe von Ordinalzahlen und des Wohlordnungssatzes, wobei $M + M = (M \times \{0\}) \cup (M \times \{1\})$ für alle Mengen M . Zunächst gilt:

Übung

Es gilt $|\lambda| = |\lambda + n|$ für jede Limesordinalzahl λ und jedes $n \in \mathbb{N}$.

[Induktion nach n oder mit Hilfe von $|\lambda + n| = |n + \lambda|$ und $n + \lambda = \lambda$.]

Satz (Additionssatz für unendliche Ordinalzahlen)

Sei α eine unendliche Ordinalzahl. Dann gilt $|\alpha + \alpha| = |\alpha|$.

Beweis

Sei $\alpha = \lambda + n$ mit λ Limes und $n \in \omega$ (es gilt $\lambda = \sup(\{\beta < \alpha \mid \beta \text{ Limes}\})$).

Wegen $|\alpha| = |\lambda|$ genügt es, $|\lambda + \lambda| = |\lambda|$ zu zeigen.

Für jede Ordinalzahl γ seien $L(\gamma)$, $N(\gamma)$ die eindeutig bestimmten Ordinalzahlen mit $\gamma = L(\gamma) + N(\gamma)$, $L(\gamma)$ Limes oder 0, $N(\gamma) \in \omega$.

Wir definieren $g(\gamma)$ für $\gamma < \lambda + \lambda$ durch:

$$g(\gamma) = \begin{cases} L(\gamma) + N(\gamma) \cdot 2, & \text{falls } \gamma < \lambda, \\ L(\delta) + N(\delta) \cdot 2 + 1, & \text{falls } \lambda \leq \gamma < \lambda + \lambda, \gamma = \lambda + \delta. \end{cases}$$

- Dann ist $g: W(\lambda + \lambda) \rightarrow W(\lambda)$ bijektiv, also $|\lambda + \lambda| = |\lambda|$.

Korollar (Additionssatz)

Sei M eine unendliche Menge. Dann gilt $|M + M| = |M|$.

Beweis

Sei $g: M \rightarrow W(\alpha)$ bijektiv (z. B. für $\alpha = |M|$). g existiert nach dem

- Wohlordnungssatz. Dann gilt $|M + M| = |\alpha + \alpha| = |\alpha| = |M|$.

Für den Multiplikationssatz brauchen wir eine Paarungsoperation Γ , die jedem Paar (α, β) von Ordinalzahlen eine möglichst kleine Ordinalzahl $\Gamma(\alpha, \beta)$

zuweist. Wir wollen erreichen, daß $\Gamma \upharpoonright W(\lambda)^2 : W(\lambda)^2 \rightarrow W(\lambda)$ bijektiv für viele Ordinalzahlen λ ist. Hierzu definieren wir eine Variante und Verallgemeinerung der Cantorschen Diagonalaufzählung von $\mathbb{N} \times \mathbb{N}$:

Definition (*die Ordnung $<$ auf den Ordinalzahlpaaren*)

Wir definieren für Ordinalzahlen $\alpha, \beta, \gamma, \delta$:

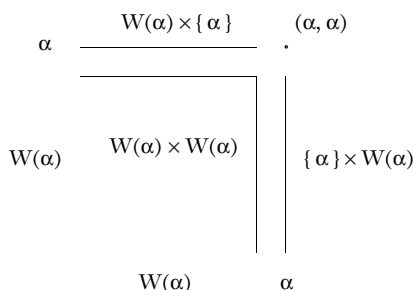
$(\alpha, \beta) < (\gamma, \delta)$, falls einer der folgenden drei Fälle eintritt:

- (i) $\max(\alpha, \beta) < \max(\gamma, \delta)$,
- (ii) $\max(\alpha, \beta) = \max(\gamma, \delta)$, $\alpha < \gamma$,
- (iii) $\max(\alpha, \beta) = \max(\gamma, \delta)$, $\alpha = \gamma$, $\beta < \delta$.

Für eine Ordinalzahl α heißt $\langle W(\alpha)^2, < \upharpoonright W(\alpha)^2 \rangle$ die *kanonische Wohlordnung* auf $W(\alpha)^2$.

Wir schreiben kurz auch $\langle W(\alpha)^2, < \rangle$ für $\langle W(\alpha)^2, < \upharpoonright W(\alpha)^2 \rangle$.

Die Ordnung $<$ kann man wie folgt visualisieren: Ist sie für $W(\alpha) \times W(\alpha)$ definiert, so entsteht die Ordnung auf $W(\alpha + 1) \times W(\alpha + 1)$ durch Addition von $W(\alpha) \times \{\alpha\}$, anschließender Addition von $\{\alpha\} \times W(\alpha)$, jeweils mit der natürlichen Ordnung, und anschließender Addition von $\{(\alpha, \alpha)\}$. Im Gegensatz zur Cantorschen Diagonalaufzählung ist diese Ordnung also aus Waagrechten



und Senkrechten zusammengesetzt, was oberhalb von ω natürlicher ist, da dann keine guten Diagonalen mehr zur Verfügung stehen. Die ersten Elemente der Ordnung sind: $(0, 0), (0, 1), (1, 0), (1, 1), (0, 2), (1, 2), (2, 0), (2, 1), (2, 2), \dots, (0, \omega), (1, \omega), (2, \omega), \dots, (\omega, 0), (\omega, 1), \dots, (\omega, \omega), (0, \omega + 1), \dots$

Da die Ordnung als Summe von Wohlordnungen definiert ist, ist sie selbst eine Wohlordnung:

Übung

$\langle W(\alpha)^2, < \rangle$ ist eine Wohlordnung für alle α . Weiter gilt: Für alle $\alpha < \beta$ ist $\langle W(\alpha)^2, < \rangle$ das durch $(0, \alpha)$ gegebene Anfangsstück von $\langle W(\beta)^2, < \rangle$.

Wir betrachten $<$ als die *kanonische Wohlordnung auf den Ordinalzahlpaaren* und definieren $W(\alpha, \beta)$ analog zu $W(\alpha)$:

Definition (*$W(\alpha, \beta)$ und die Paarungsoperation Γ*)

Seien α, β Ordinalzahlen. Wir definieren:

$$W(\alpha, \beta) = \{ (\gamma, \delta) \mid (\gamma, \delta) < (\alpha, \beta) \}, \quad \Gamma(\alpha, \beta) = \text{o.t.}(\langle W(\alpha, \beta), < \rangle).$$

Übung

Γ ist eine bijektive Operation zwischen allen Paaren von Ordinalzahlen und allen Ordinalzahlen. Für alle α ist $\Gamma \upharpoonright W(\alpha)^2 : W(\alpha)^2 \rightarrow W(\Gamma(0, \alpha))$ bijektiv.

Wir werden zeigen, daß $\Gamma \upharpoonright W(\kappa)^2 : W(\kappa)^2 \rightarrow W(\kappa)$ bijektiv ist für alle Kardinalzahlen $\kappa \geq \omega$. Diese Aussage ist nach obiger Übung äquivalent zu $\Gamma(0, \kappa) = \kappa$. Weiter ist offenbar $\Gamma(0, \alpha) \geq \alpha$ für alle α , und damit ist die Bijektivität von $\Gamma \upharpoonright W(\kappa)^2 : W(\kappa)^2 \rightarrow W(\kappa)$ äquivalent zu $\Gamma(0, \kappa) \leq \kappa$.

Satz (*Multiplikationssatz für Kardinalzahlen*)

Sei κ eine unendliche Kardinalzahl. Dann gilt $\Gamma(0, \kappa) = \kappa$, d.h.
 $\Gamma \upharpoonright W(\kappa)^2 : W(\kappa)^2 \rightarrow W(\kappa)$ ist bijektiv. Insbesondere ist $\kappa \cdot \kappa = \kappa$.

Beweis

Annahme nicht. Sei dann κ die kleinste unendliche Kardinalzahl mit $\Gamma(0, \kappa) > \kappa$. Dann existiert ein $(\alpha, \beta) \in W(0, \kappa) = W(\kappa) \times W(\kappa)$ mit $\kappa = \Gamma(\alpha, \beta)$.

κ ist als unendliche Kardinalzahl eine Limesordinalzahl und wegen $\Gamma(0, \omega) = \omega$ gilt $\kappa > \omega$. Wegen $\alpha, \beta < \kappa$ gibt es also ein $\gamma \geq \omega$ mit $\alpha, \beta < \gamma < \kappa$. Dann gilt aber

$$\kappa = \Gamma(\alpha, \beta) < \Gamma(0, \gamma).$$

Also gilt

$$|\gamma| \leq \gamma < \kappa < \Gamma(0, \gamma),$$

und damit ist $|\gamma| < \kappa = |\kappa| \leq |\Gamma(0, \gamma)| = |W(\gamma)^2| = |W(|\gamma|)^2|$, also

– $\Gamma(0, |\gamma|) > |\gamma|$, im Widerspruch zu $|\gamma| \geq \omega$ und der minimalen Wahl von κ .

Korollar (*Multiplikationssatz*)

Sei M eine unendliche Menge. Dann gilt $|M \times M| = |M|$.

Beweis

– Sei $\kappa = |M|$. Dann gilt $|M \times M| = |\kappa \cdot \kappa| = \kappa = |M|$.

Wir haben also einen zweiten (und sehr klaren) Beweis des Multiplikationssatzes gefunden. Darüber hinaus sind wir nun im Besitz einer Paarungsfunktion auf allen Ordinalzahlen.

Der obige Beweis geht auf Harward (1905), Jourdain (1908) und Hausdorff (1914) zurück, wobei dort etwas andere Paarungsfunktionen konstruiert werden: Jourdain vergleicht Paare (α, β) und (γ, δ) über die Summen $\alpha + \beta$ und $\gamma + \delta$, Harward und Hausdorff ordnen nur Paare (α, β) mit $\alpha < \beta$, was „die Hälfte“ der Ordnung $<$ liefert. Hausdorff sieht seinen Beweis als Vereinfachung von Jourdain's Beweis an, und gelangt so unwissend zu dem bereits von Harward im Anhang seiner Arbeit von 1905 skizzierten Argument. Ironischerweise war Jourdain diese Arbeit von Harward bekannt. (Siehe [Deiser 2005] für eine Darstellung der Geschichte des Multiplikationssatzes.)

Die Operation Γ wird zuweilen auch als Gödelsche Paarungsfunktion bezeichnet (vgl. [Gödel 1940]). Die Ordnung $<$ erhält man aus der Ordnung bei Harward und Hausdorff, indem die Paare (α, β) mit $\beta \leq \alpha$ wie oben mit eingewebt werden. Die zwei Beweise des Multiplikationssatzes von Hessenberg (1906, 1907) besprechen wir unten bei der Cantorsche Normalform und der direkten Definition der Exponentiation.

Die Cantorsche Normalform

Die Ordinalzahlpotenzen ω^α haben eine Reihe interessanter Eigenschaften, die wir im folgenden näher untersuchen wollen.

Definition (Hauptzahlen)

Eine Ordinalzahl β heißt eine (*additive*) *Hauptzahl*, falls es eine Ordinalzahl α gibt mit $\beta = \omega^\alpha$.

Zunächst betrachten wir für eine Ordinalzahl $\gamma \geq 1$ die Stelle α etwas genauer, für die der Schritt von ω^α zu $\omega^{\alpha+1}$ die Zahl γ erfaßt:

Definition (ω^α -Ausschöpfung)

Sei $\gamma \geq 1$ eine Ordinalzahl. Wir setzen:

$E(\gamma) =$ „die größte Ordinalzahl α mit $\omega^\alpha \leq \gamma$ “,

$N(\gamma) =$ „das größte $n \in \mathbb{N}$ mit $\omega^{E(\gamma)} \cdot n \leq \gamma$ “,

$R(\gamma) =$ „die eindeutige Ordinalzahl ρ mit $\omega^{E(\gamma)} \cdot N(\gamma) + \rho = \gamma$ “.

Dies ist wohldefiniert, und es gilt $0 \leq R(\gamma) < \omega^{E(\gamma)} \leq \gamma$.

Große natürliche Zahlen kann man z. B. in Dezimalschreibweise „exponentiell verkürzt“ angeben. So ist etwa $1304 = 10^3 \cdot 1 + 10^2 \cdot 3 + 10^0 \cdot 4$. Die Cantorsche Normalform ist nun ein Analogon für Ordinalzahlen, wobei ω statt 10 (oder allgemeiner statt einem $n \in \mathbb{N}$, $n \geq 2$) als Basis gewählt wird.

Satz (Cantorsche Normalform zur Basis ω)

Jede Ordinalzahl γ läßt sich in eindeutiger Weise schreiben als

$$\gamma = \omega^{\alpha_1} \cdot n_1 + \dots + \omega^{\alpha_k} \cdot n_k,$$

wobei $k \in \mathbb{N}$, $1 \leq n_i < \omega$ für $1 \leq i \leq k$, $\alpha_1 > \alpha_2 > \dots > \alpha_k \geq 0$.

(Der Fall $k = 0$ entspricht $\gamma = 0$.)

Beweis

Wir zeigen die Existenz und Eindeutigkeit der Normalform durch Induktion über alle Ordinalzahlen γ .

Für $\gamma = 0$ ist nichts zu zeigen.

Sei also $\gamma > 0$. Wir setzen $\alpha_1 = E(\gamma)$, $n_1 = N(\gamma)$.

Wegen $R(\gamma) < \gamma$ existiert nach Induktionsvoraussetzung eine Normalform von $R(\gamma)$, die wir in der Form

$$R(\gamma) = \omega^{\alpha_2} \cdot n_2 + \dots + \omega^{\alpha_k} \cdot n_k$$

notieren können. Aber es gilt $\gamma = \omega^{\alpha_1} \cdot n_1 + R(\gamma)$, und damit ist wegen $\omega^{\alpha_2} \leq R(\gamma) < \omega^{\alpha_1}$ offenbar eine Normalform von γ gefunden.

Ist nun $\gamma = \omega^{\alpha_1} \cdot n_1 + \dots + \omega^{\alpha_k} \cdot n_k$ irgendeine Normaldarstellung von γ , so ist $\omega^{\alpha_2} \cdot n_2 + \dots + \omega^{\alpha_k} \cdot n_k \leq \omega^{\alpha_2} \cdot (n_2 + \dots + n_k) < \omega^{\alpha_2} \cdot \omega \leq \omega^{\alpha_1}$.

Hieraus folgt, daß $\alpha_1 = E(\gamma)$ und $n_1 = N(\gamma)$ gelten muß.

Nach Induktionsvoraussetzung ist die Normaldarstellung

– von $R(\gamma)$ eindeutig, und damit ist auch die von γ selbst eindeutig.

Cantor hat die Normaldarstellung in [Cantor 1897, § 19] nur für abzählbare Ordinalzahlen aufgestellt. Der Beweis von Cantor bleibt aber für alle Ordinalzahlen α richtig.

Es gibt auch eine Normalform zur Basis 2. Sie lautet:

Übung

Jede Ordinalzahl γ läßt sich in eindeutiger Weise schreiben als

$$\gamma = 2^{\alpha_1} + \dots + 2^{\alpha_k}, \text{ wobei } k \in \mathbb{N}, \alpha_1 > \alpha_2 > \dots > \alpha_k \geq 0.$$

Allgemeiner existiert eine eindeutige Normalform zu jeder Basis $\beta \geq 2$. Die Terme der Summe sind dann von der Form $\beta^\alpha \cdot \delta$, mit $1 \leq \delta < \beta$. Zwei Zahlen γ_1 und γ_2 in Normalform zur Basis β können der Größe nach dann leicht verglichen werden, indem von „links nach rechts“ Exponenten und Koeffizienten auf einen Unterschied untersucht werden (völlig analog zu $99 < 100$ und $1046 < 1052$).

Die Hauptzahlen lassen sich nun leicht charakterisieren:

Übung

Sei γ eine unendliche Ordinalzahl. Dann sind äquivalent:

- (i) γ ist eine Hauptzahl.
- (ii) γ ist abgeschlossen unter Addition, d. h. $\delta_1 + \delta_2 < \gamma$ für alle $\delta_1, \delta_2 < \gamma$.
- (iii) Es gilt $\beta + \gamma = \gamma$ für alle $\beta < \gamma$.

Es gibt eine analoge Charakterisierung für die Ordinalzahlmultiplikation, und hierbei spielt zudem die Paarungsoperation Γ eine Rolle:

Übung

Sei γ eine unendliche Ordinalzahl. Dann sind äquivalent:

- (i) γ ist eine multiplikative Hauptzahl, d. h. es gibt eine Ordinalzahl α mit $\gamma = \omega^{\omega^\alpha}$.
- (ii) γ ist abgeschlossen unter Multiplikation, d. h. es gilt $\delta_1 \cdot \delta_2 < \gamma$ für alle $\delta_1, \delta_2 < \gamma$.
- (iii) Es gilt $\beta \cdot \gamma = \gamma$ für alle $1 \leq \beta < \gamma$.
- (iv) Es gilt $\gamma = \Gamma(0, \gamma)$ für die kanonische Paarungsfunktion Γ .

Für alle multiplikativen Hauptzahlen γ ist also die Paarungsfunktion Γ eine einfach definierte Bijektion von $W(\gamma)^2$ nach $W(\gamma)$.

Eine wichtige Anwendung der Normalform ist eine auf den ersten Blick etwas kauzige Variante der Ordinalzahladdition.

Definition (*Hessenbergsumme und Paarungssumme*)

Seien γ, δ Ordinalzahlen, und sei

$$\gamma = \omega^{\alpha_1} \cdot n_1 + \dots + \omega^{\alpha_k} \cdot n_k,$$

$$\delta = \omega^{\alpha_1} \cdot m_1 + \dots + \omega^{\alpha_k} \cdot m_k$$

die *gemeinsame Normaldarstellung* von γ und δ , d. h. wir lassen $n_i = 0$ oder $m_i = 0$ zu, um eine identische absteigende Folge von Exponenten zu erhalten; jedoch soll gelten $n_i \neq 0$ oder $m_i \neq 0$.

Dann ist die *Hessenbergsumme* von γ und δ , in Zeichen $\gamma \# \delta$, und die *Paarungssumme* von γ und δ , in Zeichen $\gamma \& \delta$, definiert durch:

$$\gamma \# \delta = \omega^{\alpha_1} \cdot (n_1 + m_1) + \dots + \omega^{\alpha_k} \cdot (n_k + m_k),$$

$$\gamma \& \delta = \omega^{\alpha_1} \cdot \pi(n_1, m_1) + \dots + \omega^{\alpha_k} \cdot \pi(n_k, m_k),$$

wobei $\pi : W(\omega) \times W(\omega) \rightarrow W(\omega)$ die Cantorsche Paarungsfunktion ist.

Satz

Sei α eine Hauptzahl. Dann ist $\& : W(\alpha) \times W(\alpha) \rightarrow W(\alpha)$ bijektiv.

Beweis

Wir haben $\gamma \& \delta < \alpha$ für alle $\gamma, \delta < \alpha$, denn α ist als Hauptzahl abgeschlossen unter Addition, und damit auch unter Paarungssummen.

Bijektivität folgt aus der Bijektivität der Cantorschen Paarungsfunktion

- π und der Eindeutigkeit der Normalform (und aus $\pi(0, 0) = 0$).

Korollar (*Multiplikationssatz*)

Sei M eine unendliche Menge. Dann gilt $|M \times M| = |M|$.

Beweis

Sei $\alpha = |M|$. Dann ist α eine unendliche Kardinalzahl und damit abgeschlossen unter Addition und somit eine Hauptzahl.

- Dann aber $|M| = |W(\alpha)| = |W(\alpha) \times W(\alpha)| = |M \times M|$.

In dieser Weise hat Hessenberg 1906 den Multiplikationssatz bewiesen. Er verwendet allerdings die Funktion $\#$ statt $\&$, und braucht dann ein zusätzliches Argument, da die Hessenbergsumme nicht injektiv ist. Sie ist aber fast injektiv: Für jedes α existieren nur endlich viele Paare (γ, δ) mit $\gamma \# \delta = \alpha$ (!). Dies genügt Hessenberg für einen Beweis. Die naheliegende Bemühung der Cantorfunktion macht diesen Beweis des Multiplikationssatzes, der fast nur aus Definitionen besteht, aber noch etwas transparenter.

Die Hessenbergsumme $\#$ ist darüber hinaus in der Beweistheorie von Interesse. In der Mengenlehre haben wir die folgende kombinatorische Eigenschaft:

Übung

Seien $\alpha_1, \dots, \alpha_n$ Ordinalzahlen. Dann gilt:

$$\alpha_1 \# \dots \# \alpha_n = \sup \{ \text{o.t.} (\bigcup_{1 \leq i \leq n} A_i) \mid A_i \text{ ist eine Menge von Ordinalzahlen mit o.t.}(A_i) = \alpha_i \text{ für } 1 \leq i \leq n \}.$$

Eine direkte Definition der Exponentiation von Ordinalzahlen

Die Exponentiation von Ordinalzahlen haben wir durch iterierte Anwendung der Multiplikation definiert. Auch die oben direkt definierten Operationen der Multiplikation und der Addition kann man rekursiv aus jeweils einfacheren Operationen gewinnen.

Übung

Sei α eine Ordinalzahl. Wir definieren $\alpha \cdot \beta$ durch Rekursion über alle Ordinalzahlen β wie folgt:

$$\alpha \cdot 0 = 0,$$

$$\alpha \cdot (\beta + 1) = \alpha \cdot \beta + \alpha \quad \text{für } \beta \text{ Ordinalzahl,}$$

$$\alpha \cdot \lambda = \sup_{\beta < \lambda} \alpha \cdot \beta \quad \text{für } \lambda \text{ Limesordinalzahl.}$$

Zeigen Sie, daß diese Multiplikation mit der alten übereinstimmt.

Die Addition kann man rekursiv mit Hilfe der Nachfolgeroperation definieren:

$$\alpha + 0 = \alpha,$$

$$\alpha + (\beta + 1) = (\alpha + \beta) + 1 \quad \text{für } \beta \text{ Ordinalzahl,}$$

$$\alpha + \lambda = \sup_{\beta < \lambda} \alpha + \beta \quad \text{für } \lambda \text{ Limesordinalzahl.}$$

Diese Addition stimmt wie im Falle der Multiplikation mit der alten Addition überein.

Umgekehrt stellt sich nun die Frage, ob die Exponentiation von Ordinalzahlen direkt, ohne ein rekursives Vorgehen, definiert werden kann. Es ist nicht so ohne weiteres klar, welche Menge als ein natürlicher Träger einer solchen direkten Exponentiation in Frage kommt.

Die Antwort gibt die folgende Konstruktion.

Definition (die natürliche Exponentiation zweier Wohlordnungen)

Seien $\langle M, < \rangle$ und $\langle N, < \rangle$ Wohlordnungen.

Dann ist die *natürliche Exponentiation* $\langle E, <^* \rangle$ von $\langle M, < \rangle$ und $\langle N, < \rangle$, in Zeichen

$$\langle E, <^* \rangle = \langle M, < \rangle^{\langle N, < \rangle},$$

wie folgt definiert. Sei

$$E = \{ f \in {}^N M \mid f(x) \neq 0_M \text{ für nur endlich viele } x \in N \},$$

wobei 0_M das Anfangselement von M sei.

(Also $E = \emptyset$, falls $M = \emptyset$, $N \neq \emptyset$ und $E = \{ \emptyset \}$, falls $N = \emptyset$.)

Für $f, g \in E$ mit $f \neq g$ sei $\Delta(f, g) = \max(\{ x \in N \mid f(x) \neq g(x) \})$.

Wir setzen dann für $f, g \in E$:

$$f <^* g \text{ falls } f \neq g \text{ und } f(\Delta) < g(\Delta) \text{ für } \Delta = \Delta(f, g).$$

$<^*$ heißt die *natürliche Ordnung* auf E .

Übung

Seien $\langle M, < \rangle, \langle N, < \rangle$ Wohlordnungen. Dann ist $\langle E, <^* \rangle = \langle M, < \rangle^{\langle N, < \rangle}$ eine Wohlordnung.

Tatsächlich entspricht nun die natürliche Exponentiation von Wohlordnungen der rekursiv definierten Exponentiation der Ordinalzahlen [Hausdorff 1906b, Hessenberg 1907]:

Satz (*Hausdorff-Hessenberg Darstellung von α^β*)

Seien $\langle M, < \rangle$ und $\langle N, < \rangle$ Wohlordnungen, und sei $\langle E, <^* \rangle = \langle M, < \rangle^{\langle N, < \rangle}$.

Sei $\alpha = \text{o.t.}(\langle M, < \rangle)$, $\beta = \text{o.t.}(\langle N, < \rangle)$. Dann gilt

$\alpha^\beta = \text{o.t.}(\langle E, <^* \rangle)$.

Insbesondere also $\alpha^\beta = \text{o.t.}(\langle W(\alpha), < \rangle^{W(\beta), <})$.

Beweis

Durch Induktion nach β bei festem α .

- Die einfachen Details seien dem Leser zur Übung überlassen.

Ein bemerkenswertes Resultat! Die rekursiv definierte Exponentiation α^β kann durch eine Ordnung dargestellt werden, deren Träger eine kleine Teilmenge von ${}^N M$ ist, wobei $\text{o.t.}(\langle M, < \rangle) = \alpha$, $\text{o.t.}(\langle N, < \rangle) = \beta$. Der Träger E besteht aus denjenigen Funktionen von N nach M , die nur an endlich vielen Stellen nichttriviale Werte annehmen. Die Ordnung $<^*$ ist über die *maximale* Stelle $\Delta(f, g)$ eines Unterschiedes von zwei Funktionen definiert (eine Definition über die *minimale* Stelle $\delta(f, g)$ eines Unterschiedes ergäbe im Normalfall keine Wohlordnung mehr: betrachte wieder $M = \{0, 1\}$, $N = \mathbb{N}$). Die Ordnung $<^*$ ist also die *antilexikographische* Ordnung auf der Teilmenge E von ${}^N M$.

Wir haben die natürliche Exponentiation hier *ad hoc* präsentiert. Sie läßt sich analytisch durch eine genaue Betrachtung des Limeschrittes der rekursiv definierten Exponentiation gewinnen. Der Leser ist aufgefordert, dies für den Fall ω^ω zu tun. Dabei ist die Beobachtung hilfreich, daß die Potenz ω^ω einer unendlichen Multiplikation nach links entspricht, also $\omega^\omega = \dots \cdot \omega \cdot \dots \cdot \omega \cdot \omega$. Dies erklärt auch das Auftreten der maximalen Stellen Δ eines Unterschiedes in obiger Definition.

Aus der direkten Definition der Ordinalzahlexponentiation erhalten wir einen weiteren Beweis des Multiplikationssatzes. Zunächst folgt aus dem Satz oben, daß die Größe von α^β nur von der Größe der Basis und des Exponenten abhängt:

Korollar

Seien α, β Ordinalzahlen. Dann ist $|\alpha^\beta| = | |\alpha|^{|\beta|} |$
(mit Ordinalzahlexponentiationen α^β und $|\alpha|^{|\beta|}$).

Beweis

Seien $\langle E_1, <^* \rangle = \langle W(\alpha), < \rangle^{W(\beta), <}$, $\langle E_2, <^* \rangle = \langle W(|\alpha|), < \rangle^{W(|\beta|), <}$.

Dann ist $|E_1| = |E_2|$ (!). Aber $\text{o.t.}(\langle E_1, <^* \rangle) = \alpha^\beta$ und $\text{o.t.}(\langle E_2, <^* \rangle) = |\alpha|^{|\beta|}$,

- also $|\alpha^\beta| = |E_1| = |E_2| = | |\alpha|^{|\beta|} |$.

Zum Vergleich: $|\alpha + \alpha| = | |\alpha| + |\alpha| |$ und $|\alpha \cdot \alpha| = | |\alpha| \cdot |\alpha| |$ sind leicht zu sehen, da Addition und Multiplikation über Vereinigung und Kreuzprodukt direkt definiert sind. Aus der rekursiven Definition für die Exponentiation ist dagegen $|2^\alpha| = |2^{|\alpha|}|$ keineswegs klar. Die Aussage läßt sich induktiv beweisen, wenn der Multiplikationssatz benutzt wird, was wir hier gerade vermeiden wollen. Die Hausdorff-Hessenberg-Darstellung von 2^α liefert direkt, daß die Größe von 2^α nur von der Größe des Exponenten abhängt.

Zur Vorbereitung des Arguments betrachten wir folgenden Beweis für die Abgeschlossenheit einer unendlichen Kardinalzahl unter Addition.

Satz ($2 \cdot \kappa = \kappa$ und additive Abgeschlossenheit unendlicher Kardinalzahlen)

Sei κ eine unendliche Kardinalzahl. Dann gilt $2 \cdot \kappa = \kappa$.

Insbesondere ist $\alpha + \beta < \kappa$ für alle $\alpha, \beta < \kappa$.

Beweis

Annahme nicht. Sei dann o. E. κ das kleinste Gegenbeispiel. Wegen

$\kappa < 2 \cdot \kappa = \sup \{ 2 \cdot \alpha \mid \alpha < \kappa \}$ existiert dann ein $\alpha < \kappa$ mit $\kappa \leq 2 \cdot \alpha$.

Dann ist $|2 \cdot |\alpha|| = |2 \cdot \alpha| \geq \kappa > \alpha \geq |\alpha|$ und $\omega \leq |\alpha| < \kappa$, im Widerspruch zur Minimalität von κ .

Zusatz: Für $\alpha \leq \beta < \kappa$ ist $|\alpha + \beta| \leq |\beta + \beta| = |\beta \cdot 2| = |2 \cdot \beta| \leq 2 \cdot \beta < \kappa$,

– also auch $\alpha + \beta < \kappa$, da κ Kardinalzahl.

Ein weiteres Argument ist: Es gilt allgemein $2 \cdot \lambda = \lambda$ für alle Limesordinalzahlen λ , wie eine Induktion nach λ zeigt. Jede Kardinalzahl $\kappa \geq \omega$ ist aber eine Limesordinalzahl.

Analog folgt nun:

Satz ($2^\kappa = \kappa$ und multiplikative Abgeschlossenheit unendlicher Kardinalzahlen)

Sei κ eine unendliche Kardinalzahl. Dann ist $2^\kappa = \kappa$ (ordinalzahlarithmetisch).

Insbesondere ist $\alpha \cdot \beta < \kappa$ für alle $\alpha, \beta < \kappa$.

Beweis

Annahme nicht. Sei dann o. E. κ das kleinste Gegenbeispiel.

Wegen $\kappa < 2^\kappa = \sup \{ 2^\alpha \mid \alpha < \kappa \}$ existiert ein $\alpha < \kappa$ mit $2^\alpha \geq \kappa$. Dann ist aber $|2^{|\alpha|}| = |2^\alpha| \geq \kappa > |\alpha| \geq \omega$, im Widerspruch zur Minimalität von κ .

Zusatz: Sind $\alpha \leq \beta < \kappa$, so ist $\beta + \beta < \kappa$ und damit:

– $\alpha \cdot \beta \leq \beta \cdot \beta \leq 2^\beta \cdot 2^\beta = 2^{\beta+\beta} < 2^\kappa = \kappa$.

Korollar (Multiplikationssatz)

Sei M eine unendliche Menge. Dann ist $|M \times M| = |M|$.

Beweis

Sei $\kappa = |M|$. Dann ist $|M \times M| = |\kappa \cdot \kappa| < \kappa^+$ wegen $\kappa \cdot \kappa < \kappa^+$.

– Hieraus folgt $|M \times M| = \kappa$.

Dieser Beweis des Multiplikationssatzes folgt [Hessenberg 1907]. Das Argument ist eine hübsche Anwendung der direkten Definition der Exponentiation.

Schließlich halten wir fest:

Korollar (*Exponentiationssatz*)

- (i) Seien α, β Ordinalzahlen mit $\alpha^\beta \geq \omega$.
Dann gilt $|\alpha^\beta| = \max(|\alpha|, |\beta|)$.
- (ii) Seien $\langle M, < \rangle, \langle N, < \rangle$ Wohlordnungen, und sei $\langle E, <^* \rangle = \langle M, < \rangle^{\langle N, < \rangle}$.
Ist E unendlich, so gilt $|E| = \max(|M|, |N|)$.

Beweis

zu (i): Sei $\kappa = \max(|\alpha|, |\beta|)$. Dann ist $\kappa \geq \omega$.

Offenbar $\kappa \leq \alpha^\beta$, also auch $\kappa \leq |\alpha^\beta|$ wegen κ Kardinalzahl.

Es gilt $\alpha^\beta \leq (2^\alpha)^\beta = 2^{\alpha \cdot \beta}$. Also $|\alpha^\beta| \leq |2^{\alpha \cdot \beta}| = |2^\kappa| = \kappa$.

zu (ii): Sei $\alpha = \text{o.t.}(\langle M, < \rangle)$, $\beta = \text{o.t.}(\langle N, < \rangle)$.

— Dann gilt $|E| = |\alpha^\beta|$, und die Behauptung folgt also aus (i).

Die systematische Untersuchung linearer Ordnungen war einer der Schwerpunkte der mengentheoretischen Forschung von Felix Hausdorff. Ziel war ein systematischer Aufbau aller Typen von linearen Ordnungen durch arithmetische Operationen. Hausdorffs Vision einer „Beherrschung der Typenwelt“ erfüllte sich nur zum Teil. Die Wege von einfachen zu komplizierteren linearen Ordnungen erwiesen sich als zu verzweigt, um durch eine Strukturtheorie vollständig vermessen und kartografiert werden zu können. Man wird also mit Teilresultaten zufrieden sein müssen, zumal die Theorie der Wohlordnungen schon komplex genug ist. Eine Methode, die Hausdorff in der Verhaltensforschung linearer Ordnungen einsetzte, werden wir im nächsten Kapitel noch kennenlernen.

**Aus der Einleitung von Felix Hausdorffs
„Grundzüge einer Theorie der geordneten Mengen“**

„Im folgenden wird zum ersten Male eine systematische und möglichst allgemein gehaltene Einführung in das von Herrn G. Cantor erschlossene, aber noch so gut wie unbekannte Gebiet der einfach [linear] geordneten Mengen versucht, von denen bisher eigentlich nur die wohlgeordneten Mengen und die Mengen reeller Zahlen eine eingehende Betrachtung erfahren haben. Auch meine eignen Vorarbeiten, von denen hier übrigens nichts vorausgesetzt wird, verfolgen im ganzen eine speziellere Richtung ... In der vorliegenden Abhandlung, sollte sie nicht zum Buche anschwellen, mußten diese spezielleren [Ordnungs-] Typen durchaus in die Rolle gelegentlicher Illustrationsbeispiele zurücktreten und die allgemeinen Methoden den Vordergrund einnehmen. Das vorschwebende Ideal war etwa eine Beherrschung der Typenwelt in dem Sinne, daß die komplexen Gebilde durch erzeugende Operationen aus elementaren aufgebaut erscheinen sollten ...“

(Felix Hausdorff 1908)

9. Große Teilmengen und große Kardinalzahlen

Wir verlassen nun die Arithmetik von linearen Ordnungen, und untersuchen spezielle Eigenschaften, die einzelne lineare Ordnungstypen haben können, etwas genauer. Das hierbei zentrale Konzept der Konfinalität führt uns direkt zu den 1908 von Hausdorff entdeckten unerreichbaren Kardinalzahlen. Während Hausdorff vor ihrer Größe zurückschreckte, untersuchte Paul Mahlo (1883 – 1971) zwischen 1911 und 1913 ganze Hierarchien von unerreichbaren Kardinalzahlen, und fand dabei Zahlen, die alle diese Hierarchien transzendieren. Die dabei auftretenden Begriffe ergeben sich aus heutiger Sicht am natürlichsten aus einer Mathematisierung der Idee „große Teilmenge von“. Die sich hieraus ergebende Technologie liefert dann – ahistorisch, aber zwanglos –, nicht nur eine Flugverbindung zu den Mahlo-Kardinalzahlen, sondern startet darüber hinaus noch das Weltraumprogramm der meßbaren Kardinalzahlen. Die ersten Lebenszeichen der meßbaren Kardinalzahlen hatte Stanisław Ulam (1909 – 1984) im Jahre 1930, ihrer in der Maßtheorie spürbaren Hintergrundstrahlung nachgehend, auf ganz anderem Wege gefunden und untersucht.

Konfinalitäten

Wir hatten schon das Phänomen beobachtet, daß die relativ große Kardinalzahl $\aleph_\omega$ das Supremum der relativ kleinen Menge $\{\aleph_n \mid n \in \omega\}$ ist. Allgemein kann man für lineare Ordnungen die Größe von unbeschränkten Teilordnungen betrachten. Dies führt zu einem fundamentalen Begriff der modernen Mengenlehre, der auf Hessenberg und Hausdorff zurückgeht:

Definition (*Konfinalität einer linearen Ordnung, $cf(\alpha)$*)

Sei $\langle M, < \rangle$ eine lineare Ordnung, und sei $N \subseteq M$.

N heißt *konfinal in* $\langle M, < \rangle$, falls für alle $x \in M$ ein $y \in N$ existiert mit $x \leq y$.

Die *Konfinalität von* $\langle M, < \rangle$, in Zeichen $cf(\langle M, < \rangle)$, ist definiert durch:

$$cf(\langle M, < \rangle) = \min \{ |N| \mid N \subseteq M, N \text{ konfinal in } \langle M, < \rangle \}.$$

Wir schreiben $cf(\alpha)$ für $cf(\langle W(\alpha), < \rangle)$ für Ordinalzahlen α .

Ist $N \subseteq W(\alpha)$ konfinal in $\langle W(\alpha), < \rangle$, so sagen wir kurz: N ist konfinal in α .

cf steht für engl. *cofinality*. Hausdorff hat den Term eingeführt und den Begriff untersucht. Analog definiert er: $N \subseteq M$ heißt *koinitial* in $\langle M, < \rangle$, falls für alle $x \in M$ ein $y \in N$ existiert mit $y \leq x$.

tiert mit $y \leq x$. Man kann dann die *Koinitialität* $ci(\langle M, < \rangle)$ von $\langle M, < \rangle$ definieren als die kleinste Kardinalität einer koinitalen Teilmenge von M . Es gilt $ci(\langle M, < \rangle) = cf(\langle M, <^* \rangle)$ mit $<^* = <^{-1}$. Für jeden Schnitt (L, R) von $\langle M, < \rangle$ (vgl. Kapitel 10) sind dann $cf(\langle L, < \rangle)$ und $ci(\langle R, < \rangle)$ wichtige Größen zur Beschreibung des Verhaltens der linearen Ordnung an der „Stelle“ (L, R) .

Es gilt $cf(0) = cf(\langle \emptyset, \emptyset \rangle) = 0$. Hat $\langle M, < \rangle$ ein größtes Element x , so ist $\{x\}$ konfinal in der Ordnung, also $cf(\langle M, < \rangle) = 1$. Andernfalls gilt $cf(\langle M, < \rangle) \geq \omega$. Speziell gilt also $cf(\alpha) = 1$ für alle Nachfolgerordinalzahlen α und $cf(\lambda) \geq \omega$ für alle Limesordinalzahlen λ . Trivialerweise gilt $cf(\alpha) \leq \alpha$ für alle α .

Definition (*regulär und singular*)

Sei α eine Ordinalzahl. α heißt *regulär*, falls $cf(\alpha) = \alpha$.

Andernfalls heißt α *singular*.

Offenbar sind 0, 1 und ω regulär, 2, 3, 4, $\omega + \omega$, $\omega \cdot \omega$ sind singular. Regularität kann nur den Kardinalzahlen zukommen.

Übung

(i) Ist eine Ordinalzahl α regulär, so ist α eine Kardinalzahl.

(ii) Es gilt $cf(\aleph_\lambda) = cf(\lambda)$ für alle Limesordinalzahlen λ .

[Betrachte $N = \{ \aleph_\alpha \mid \alpha < \lambda \} \subseteq W(\aleph_\lambda)$.]

Insbesondere gilt $cf(\aleph_\omega) = cf(\aleph_{\omega_0}) = cf(\aleph_{\omega_1 + \omega}) = \omega$, $cf(\aleph_{\omega_1}) = cf(\omega_1)$.

Die Konfinalitäten von unendlichen Nachfolgerkardinalzahlen lassen sich mit Hilfe des Multiplikationssatzes leicht bestimmen:

Satz (*Regularität von unendlichen Nachfolgerkardinalzahlen*)

$\aleph_{\alpha+1}$ ist regulär für alle Ordinalzahlen α .

Beweis

Sei $N \subseteq W(\aleph_{\alpha+1})$, $|N| \leq \aleph_\alpha$. Dann ist $|W(\gamma)| \leq \aleph_\alpha$ für alle $\gamma \in N$ und somit:

$$|\bigcup_{\gamma \in N} W(\gamma)| \leq |N| \cdot \sup_{\gamma \in N} |W(\gamma)| \leq |N| \cdot \aleph_\alpha \leq \aleph_\alpha \cdot \aleph_\alpha = \aleph_\alpha.$$

– Also ist $\bigcup_{\gamma \in N} W(\gamma) \subset W(\aleph_{\alpha+1})$, und somit ist N nicht konfinal in $\aleph_{\alpha+1}$.

Übung

Zeigen Sie den Multiplikationssatz unter der Annahme der Regularität von unendlichen Nachfolgerkardinalzahlen.

[Sei $f: M \times M \rightarrow W(\kappa^+)$ mit $\kappa = |M| \geq \aleph_0$. Für $x \in M$ sei

$\gamma_x = \sup(f''(\{x\} \times M))$. Wegen Regularität von κ^+ ist $\gamma_x < \kappa^+$ für alle $x \in M$.

Wieder wegen Regularität von κ^+ ist dann $\{\gamma_x \mid x \in M\}$ beschränkt in $W(\kappa^+)$, also ist f nicht surjektiv. Also $|M \times M| < \kappa^+$.]

Hausdorff hat bereits in [Hausdorff 1904] die Regularität von $\aleph_{\alpha+1}$ behauptet, ohne allerdings einen Beweis anzugeben. Die Regularität von Nachfolgerkardinalzahlen wird für den Beweis der in der Arbeit angegebenen „Hausdorff-Formel“ gebraucht:

Korollar (*Hausdorff-Formel*)

Seien α, β Ordinalzahlen. Dann gilt:

$$\aleph_{\alpha+1}^{\aleph_\beta} = \aleph_{\alpha+1} \aleph_\alpha^{\aleph_\beta}.$$

Beweis

Nach dem Multiplikationssatz gilt offenbar $\aleph_{\alpha+1}^{\aleph_\beta} \geq \aleph_{\alpha+1} \aleph_\alpha^{\aleph_\beta}$.
Wir zeigen noch $\leq$.

1. Fall: $\alpha + 1 \leq \beta$

$$\text{Dann ist } \aleph_{\alpha+1}^{\aleph_\beta} \leq (2^{\aleph_{\alpha+1}})^{\aleph_\beta} = 2^{\aleph_{\alpha+1} \aleph_\beta} = 2^{\aleph_\beta} \leq \aleph_{\alpha+1} \aleph_\alpha^{\aleph_\beta}.$$

2. Fall: $\beta < \alpha + 1$

Wegen $\aleph_{\alpha+1}$ regulär ist $\text{rng}(f)$ beschränkt in $W(\aleph_{\alpha+1})$ für alle $f: W(\aleph_\beta) \rightarrow W(\aleph_{\alpha+1})$. Damit gilt dann:

$$\aleph_{\alpha+1}^{\aleph_\beta} = |^{W(\aleph_\beta)} W(\aleph_{\alpha+1})| = |\bigcup_{\gamma < \aleph_{\alpha+1}} ^{W(\aleph_\beta)} W(\gamma)| \leq \aleph_{\alpha+1} \aleph_\alpha^{\aleph_\beta},$$

— denn für alle $\gamma < \aleph_{\alpha+1}$ ist $|^{W(\aleph_\beta)} W(\gamma)| \leq \aleph_\alpha^{\aleph_\beta}$ wegen $|W(\gamma)| \leq \aleph_\alpha$.

Übung

Es gilt:

(i) $\aleph_1^{\aleph_0} = 2^{\aleph_0}$, und allgemein $\aleph_{n+1}^{\aleph_n} = 2^{\aleph_n}$ für $n < \omega$.

(ii) $\aleph_2^{\aleph_0} = \aleph_2 \cdot 2^{\aleph_0}$, und allgemein $\aleph_{n+k}^{\aleph_n} = \aleph_{n+k} 2^{\aleph_n}$ für $n, k < \omega$.

Im ersten Abschnitt hatten wir die μ -Zerlegbarkeit einer Kardinalzahl κ definiert (1.12): κ ist μ -zerlegbar, falls $\kappa = \sum_{i \in I} \kappa_i$ ist für Kardinalzahlen $\kappa_i < \kappa$ und $|I| = \mu$. Zerlegbarkeit und Konfinalität hängen wie folgt zusammen:

Satz (*Konfinalität und Zerlegbarkeit*)

Sei $\kappa \geq \omega$ eine Kardinalzahl. Dann gilt:

$\text{cf}(\kappa) =$ „das kleinste μ mit κ ist μ -zerlegbar“.

Beweis

Sei $\mu =$ „das kleinste μ' mit κ ist μ' -zerlegbar“.

zu $\mu \leq \text{cf}(\kappa)$: Sei $N \subseteq W(\kappa)$ konfinal in κ , $|N| = \text{cf}(\kappa)$. Dann gilt:

$$\kappa = |\bigcup N| \leq \sum_{\alpha \in N} |\alpha| \leq |N| \cdot \kappa = \kappa.$$

Also ist $\langle \kappa_\alpha \mid \alpha \in N \rangle$ mit $\kappa_\alpha = |\alpha|$ eine $\text{cf}(\kappa)$ -Zerlegung von κ .

zu $\text{cf}(\kappa) \leq \mu$: Sei $\langle \kappa_\alpha \mid \alpha < \mu \rangle$ eine μ -Zerlegung von κ .

Sei $\lambda = \sup \{ \kappa_\alpha \mid \alpha < \mu \}$. Dann gilt $\kappa = \sum_{\alpha < \mu} \kappa_\alpha = \mu \cdot \lambda$. Also $\kappa \in \{ \mu, \lambda \}$.

Ist $\mu = \kappa$, so ist $\text{cf}(\kappa) \leq \kappa \leq \mu$. Andernfalls ist $\lambda = \kappa$, und dann ist

— $N = \{ \kappa_\alpha \mid \alpha < \mu \}$ konfinal in κ , also ebenfalls $\text{cf}(\kappa) \leq \mu$.

Zerlegungen bzw. Konfinalitäten sind in der Kardinalzahlarithmetik von großer Bedeutung. Aus dem Satz von König-Zermelo hatten wir im ersten Abschnitt schon die ω -Unzerlegbarkeit von 2^ω bewiesen. Allgemeiner halten wir nun fest:

Satz (*Konfinalitätsungleichungen*)

Sei $\kappa \geq \omega$ eine Kardinalzahl. Dann gilt:

- (i) $\kappa < \kappa^{\text{cf}(\kappa)}$,
- (ii) $\kappa < \text{cf}(\mu^\kappa)$ für alle Kardinalzahlen $\mu \geq 2$.

Beweis

zu (i):

Sei $\kappa = \sum_{i \in I} \kappa_i$ für Kardinalzahlen $\kappa_i < \kappa$ und eine Menge I mit $|I| = \text{cf}(\kappa)$.

Dann gilt nach dem Satz von König-Zermelo:

$$\kappa = \sum_{i \in I} \kappa_i < \prod_{i \in I} \kappa = \kappa^{|I|} = \kappa^{\text{cf}(\kappa)}.$$

zu (ii):

Sei I eine Menge mit $|I| \leq \kappa$, und seien $\kappa_i < \mu^\kappa$ Kardinalzahlen für $i \in I$.

Dann gilt nach dem Satz von König-Zermelo:

$$\sum_{i \in I} \kappa_i < \prod_{i \in I} \mu^\kappa \leq (\mu^\kappa)^\kappa = \mu^{\kappa \cdot \kappa} = \mu^\kappa.$$

— Also ist $\mu^\kappa \geq \omega$ nicht κ -zerlegbar, d. h. $\text{cf}(\mu^\kappa) > \kappa$.

Die erste Aussage folgt auch aus der zweiten, denn es gilt $\text{cf}(\kappa) < \text{cf}(\kappa^{\text{cf}(\kappa)})$ nach (ii), also ist $\kappa^{\text{cf}(\kappa)} = \kappa$ unmöglich. Also $\kappa^{\text{cf}(\kappa)} > \kappa$.

Es gilt z. B.:

$$|\mathbb{R}| = 2^{\aleph_0}, \aleph_1^{\aleph_0}, \aleph_\omega^{\aleph_0} \text{ sind verschieden von } \aleph_\omega, \aleph_{\omega_1 + \omega}, \aleph_{\varepsilon_0}, \text{ usw.},$$

$$2^{\aleph_1}, \aleph_0^{\aleph_1}, \aleph_1^{\aleph_1}, \aleph_2^{\aleph_1} \text{ sind verschieden von } \aleph_{\omega_1}, \aleph_{\omega_2 + \omega_1}, \aleph_{\omega_{\omega_1}}, \text{ usw.}$$

$$\aleph_\omega^{\aleph_0} > \aleph_\omega, \aleph_{\omega_1}^{\aleph_2} \geq \aleph_{\omega_1}^{\aleph_1} > \aleph_{\omega_1}.$$

Weiter gilt $2^{\text{cf}(\kappa)} \neq \kappa$ für alle $\kappa \geq \omega$, da sonst $\text{cf}(\kappa) < \text{cf}(2^{\text{cf}(\kappa)}) = \text{cf}(\kappa)$.

Die Operation $\kappa^{\text{cf}(\kappa)}$ für $\kappa \geq \omega$ ist neben 2^κ die wichtigste Operation, die uns von einer unendlichen Kardinalzahl zu einer größeren führt. Im allgemeinen ist die cf-Operation moderater, wie der alte Trick zeigt:

$$\kappa^{\text{cf}(\kappa)} \leq (2^\kappa)^{\text{cf}(\kappa)} = 2^{\kappa \cdot \text{cf}(\kappa)} = 2^\kappa.$$

Man kann noch mehr über den Zusammenhang von 2^κ und $\kappa^{\text{cf}(\kappa)}$ sagen. Hierzu zunächst ein Begriff, der auch in anderem Zusammenhang von Interesse ist:

Definition (*starke Limeskardinalzahlen*)

Sei λ eine Limesordinalzahl. $\aleph_\lambda$ heißt eine *starke Limeskardinalzahl* (oder ein *starker Limes*), falls gilt:

$$2^{\aleph_\alpha} < \aleph_\lambda \text{ für alle } \alpha < \lambda.$$

Übung

Ist κ ein starker Limes, so gilt $2^\kappa = \kappa^{\text{cf}(\kappa)}$.

[Seien $\kappa_i, i \in I, |I| = \text{cf}(\kappa)$, Kardinalzahlen mit $\kappa = \sum_{i \in I} \kappa_i$. Dann gilt

$$2^\kappa = 2^{\sum_{i \in I} \kappa_i} = \prod_{i \in I} 2^{\kappa_i} \leq \prod_{i \in I} \kappa = \kappa^{\text{cf}(\kappa)} \leq 2^\kappa.]$$

Ist κ regulär, so ist $\kappa^{\text{cf}(\kappa)} = \kappa^\kappa = 2^\kappa$. Sei also $\kappa > \omega$ singulär. Gilt $2^{\text{cf}(\kappa)} > \kappa$, so ist ebenfalls $\kappa^{\text{cf}(\kappa)} = 2^{\text{cf}(\kappa)}$, denn

$$2^{\text{cf}(\kappa)} \leq \kappa^{\text{cf}(\kappa)} \leq (2^{\text{cf}(\kappa)})^{\text{cf}(\kappa)} = 2^{\text{cf}(\kappa)}.$$

Sei also andernfalls $2^{\text{cf}(\kappa)} < \kappa$ (der Fall $2^{\text{cf}(\kappa)} = \kappa$ ist nach obigen Überlegungen ausgeschlossen). Gilt $2^\kappa = \kappa^+$, d. h. (GCH) gilt an der Stelle κ , so ist $\kappa^{\text{cf}(\kappa)} = \kappa^+$. Was aber, wenn $2^\kappa \geq \kappa^{++}$? Diese Überlegungen führen zur folgenden Hypothese, der in der Entwicklung der modernen Mengenlehre eine zentrale Rolle zukam:

Singuläre Kardinalzahlhypothese (SCH)

Sei $\kappa \geq \omega$ eine singuläre Kardinalzahl mit $2^{\text{cf}(\kappa)} < \kappa$.

Dann gilt $\kappa^{\text{cf}(\kappa)} = \kappa^+$.

(SCH) kann man als eine Abschwächung von (GCH) für singuläre Limeskardinalzahlen ansehen. In der Tat ist nun auch (SCH) wieder eine Aussage, die weder beweisbar noch widerlegbar ist – wobei diesmal, im Gegensatz zu (GCH) für die Konstruktion von Modellen, in denen (SCH) falsch ist, große Kardinalzahlaxiome benötigt werden (im Bereich von meßbaren Kardinalzahlen, s. u.). (SCH) ist also weder beweisbar noch widerlegbar, aber aus logischer Sicht viel einfacher in einem Modell zu realisieren als in einem Modell zu verletzen. Dies ist eine Asymmetrie, die für viele unabhängige Aussagen gilt.

Die Hypothese kann dann sogar für den ersten Kandidaten $\aleph_\omega$ verletzt sein: Es kann zum Beispiel $2^{\aleph_\omega} = \aleph_\omega^{\aleph_0} = \aleph_{\omega+2}$ und $\aleph_\omega$ ein starker Limes sein (modulo der Konsistenz großer Kardinalzahlen). Wie groß $\aleph_\omega^{\aleph_0}$ sein kann, wenn $\aleph_\omega$ ein starker Limes ist, ist ein offenes Problem. Man weiß, daß in diesem Fall $2^{\aleph_\omega} = \aleph_\omega^{\aleph_0} < \aleph_{\omega_4}$ gilt, und die 4 ist hier kein Druckfehler.

Weiter ist sogar $2^{\aleph_n} = \aleph_{n+1}$ für $n < \omega$ und $2^{\aleph_\omega} = \aleph_{\omega+2}$ modulo großer Kardinalzahlen möglich. Dagegen folgt aus „ $2^{\aleph_\alpha} = \aleph_{\alpha+1}$ für alle $\alpha < \omega_1$ “, daß auch $2^{\aleph_{\omega_1}} = \aleph_{\omega_1+1}$ gilt, d. h. (GCH) kann zum ersten Mal an der Stelle $\aleph_\omega$ verletzt sein, nicht aber zum ersten Mal an der Stelle $\aleph_{\omega_1}$.

All dies sind höchstgradig nichttriviale moderne Resultate. Sie belegen die subtilen Unterschiede zwischen regulär und singulär, und weiter zwischen $\text{cf}(\kappa) = \omega$ und $\text{cf}(\kappa) \geq \omega_1$. Sie belegen auch die beharrlich-forschende Belagerung der Kardinalzahlexponentiation, jener letztendlich wohl uneinnehmbaren Burg der Mengenlehre.

Definition (die Gimel-Funktion)

Sei $\kappa \geq \omega$ eine Kardinalzahl. Wir setzen:

$$\mathfrak{I}(\kappa) = \kappa^{\text{cf}(\kappa)}.$$

Gimel ist der dritte Buchstabe des hebräischen Alphabets.

Es gilt $\mathfrak{I}(\kappa) > \kappa$ für alle unendlichen Kardinalzahlen.

Zurück zu konfinalen Teilmengen. Sie haben eine nützliche Transitivitätseigenschaft, die dazu führt, daß die cf-Operation bereits nach einer Anwendung eine reguläre Kardinalzahl hinterläßt:

Übung

Sei $\langle M, < \rangle$ eine lineare Ordnung, und sei $N \subseteq M$ konfinal in $\langle M, < \rangle$.
 Weiter sei $P \subseteq N$ konfinal in $\langle N, < \rangle$. Dann ist P konfinal in $\langle M, < \rangle$.

Hieraus folgt leicht:

Satz (*Regularität einer Konfinalität*)

Sei $\langle M, < \rangle$ eine lineare Ordnung, und sei $\kappa = \text{cf}(\langle M, < \rangle)$.

Dann ist κ regulär.

Insbesondere gilt $\text{cf}(\text{cf}(\alpha)) = \text{cf}(\alpha)$ für alle Ordinalzahlen α .

Beweis

Sei $N \subseteq M$ konfinal in $\langle M, < \rangle$ mit $|N| = \kappa = \text{cf}(\langle M, < \rangle)$.

Sei $P \subseteq N$ konfinal in $\langle N, < \rangle$ mit $|P| = \text{cf}(\langle N, < \rangle)$.

Dann ist P konfinal in $\langle M, < \rangle$, also $|P| \geq \kappa$.

– Aber $|P| \leq |N| = \kappa$, also $|P| = \kappa$.

Eine sehr anschauliche Charakterisierung der Konfinalität einer linearen Ordnung ist:

Satz (*Konfinalität und aufsteigende Folgen*)

Sei $\langle M, < \rangle$ eine lineare Ordnung, und sei $\kappa = \text{cf}(\langle M, < \rangle)$. Weiter sei

$\mu =$ „die kleinste Ordinalzahl α , für die ein ordnungstreu
 $f : W(\alpha) \rightarrow M$ existiert mit $\text{rng}(f)$ konfinal in $\langle M, < \rangle$.“

Dann gilt $\mu = \kappa$.

Beweis

zu $\kappa \leq \mu$:

Sei $f : W(\mu) \rightarrow M$ ordnungstreu mit $\text{rng}(f)$ konfinal in $\langle M, < \rangle$.

Dann gilt $\kappa \leq |\text{rng}(f)| = |\mu| \leq \mu$.

zu $\mu \leq \kappa$:

Sei $N \subseteq M$ konfinal in $\langle M, < \rangle$ mit $|N| = \kappa$.

Sei $g : W(\kappa) \rightarrow N$ bijektiv.

Wir definieren für $\alpha < \kappa$ rekursiv:

$h(\alpha) = g(\beta)$, wobei $\beta =$ „das kleinste γ mit $g(\gamma) > h(\alpha')$ für alle $\alpha' < \alpha$.“

[β existiert, denn $|W(\alpha)| < \kappa$ für alle $\alpha < \kappa$, also ist $h''W(\alpha)$
 nicht konfinal in $\langle N, < \rangle$ für alle $\alpha < \kappa$.]

Dann ist $h : W(\kappa) \rightarrow N$ ordnungstreu und $\text{rng}(h)$ konfinal in $\langle N, < \rangle$ (!).

– Also ist $\mu \leq \kappa$.

Die Konfinalität einer linearen Ordnung ist also die minimale Anzahl von Schritten, die wir brauchen, um das Ende der linearen Ordnung zu erreichen (mit einer monoton-aufsteigenden Wanderung). Speziell für Ordinalzahlen α definieren wir hierzu noch:

Definition (*strikt aufsteigend, schwach aufsteigend, stetig und konfinal*)

Eine Folge $g = \langle \gamma_\beta \mid \beta < \delta \rangle$ von Ordinalzahlen heißt:

- (i) *strikt aufsteigend*, falls $\gamma_\beta < \gamma_{\beta'}$ für alle $\beta < \beta' < \delta$,
- (ii) *schwach aufsteigend*, falls $\gamma_\beta \leq \gamma_{\beta'}$ für alle $\beta < \beta' < \delta$,
- (iii) *stetig*, falls g schwach aufsteigend ist und $\sup \{ \gamma_\beta \mid \beta < \lambda \} = \gamma_\lambda$ für alle Limesordinalzahlen $\lambda < \delta$ gilt.
- (iv) *konfinal in α* , falls $\{ \gamma_\beta \mid \beta < \delta \}$ konfinal in α ist.

Strikt aufsteigende und stetige Ordinalzahlfunktionen heißen auch *normal* oder *Normalfunktionen*. Seit Cantors Definition der Ordinalzahlexponentiation wurden derartige Funktionen und ihre Fixpunkte betrachtet. Ausführlich untersucht wurden sie insbesondere von Hessenberg (1906) und von Oswald Veblen (1880 – 1960) in [Veblen 1908]. Hinreichend lange Normalfunktionen haben viele Fixpunkte (vgl. 2.8). Die Bezeichnung „Normalfunktion“ stammt von Hausdorff (1914).

Übung

Sei α eine Limesordinalzahl. Dann gilt:

- (i) Ist $\langle \gamma_\beta \mid \beta < \delta \rangle$ schwach aufsteigend konfinal in α , so ist $\text{cf}(\alpha) = \text{cf}(\delta)$.
- (ii) Es existiert eine strikt aufsteigende, stetige und in α konfinale Folge $\langle \gamma_\beta \mid \beta < \text{cf}(\alpha) \rangle$.

Der erste Teil zeigt, daß das Zulassen von beliebigen Verzögerungen die Konfinalität der Länge des Anstiegs zu einer Limesordinalzahl α invariant läßt.

Große Kardinalzahlen

Aus der Gleichung $\text{cf}(\aleph_\lambda) = \text{cf}(\lambda)$ für Limesordinalzahlen λ folgt, daß Regularität für überabzählbare Limeskardinalzahlen nur für sehr große Zahlen möglich ist: Ist λ eine Limesordinalzahl derart, daß $\aleph_\lambda$ regulär ist, so ist $\aleph_\lambda = \text{cf}(\aleph_\lambda) = \text{cf}(\lambda) \leq \lambda \leq \aleph_\lambda$. Also ist $\aleph_\lambda = \lambda$, d. h. λ ist ein Fixpunkt der Alephreihe. Wir wissen, daß Fixpunkte der Alephreihe existieren, aber Regularität ist für Fixpunkte der Alephreihe eine starke Forderung – i. a. ist die Konfinalität eines Fixpunktes $\aleph_\lambda = \lambda$ der Alephreihe echt kleiner als λ , und damit ist $\aleph_\lambda$ nicht regulär.

Übung

Sei λ die kleinste Ordinalzahl mit $\aleph_\lambda = \lambda$. Dann gilt $\text{cf}(\lambda) = \omega$.

[Sei $\kappa_0 = \aleph_0$, $\kappa_{n+1} = \aleph_{\kappa_n}$ für $n \in \omega$.

Dann ist $\lambda = \sup_{n \in \omega} \kappa_n$ und $\{ \kappa_n \mid n \in \omega \}$ ist konfinal in λ .]

Die Frage nach der Existenz von überabzählbaren regulären Limeskardinalzahlen ist ein „Sesam, öffne dich!“ am Eingangsportaal zum Reich der großen Kardinalzahlen. Wir definieren:

Definition (*schwach unerreichte Kardinalzahlen*)

Eine Limeskardinalzahl $\kappa > \omega$ heißt *schwach unerreichbar*, falls κ regulär ist.

Die schwach unerreichten Kardinalzahlen sind also genau die Limeskardinalzahlen $\aleph_\lambda$ mit $\text{cf}(\lambda) = \lambda = \aleph_\lambda$. Hausdorff stieß in seiner Untersuchung des Konfinalitätsbegriffs 1908 auf diese Zahlen:

Hausdorff (1908): „Satz III. Jede Anfangszahl [jedes ω_ξ], deren Index keine Limeszahl ist, ist regulär...

Die Frage, ob der Satz III auch umkehrbar ist oder ob es auch reguläre Anfangszahlen mit Limesindex gibt, muß hier unentschieden bleiben... Die Existenz einer solchen Zahl ξ [ξ regulär, $\xi = \omega_\xi$] erscheint hiernach mindestens problematisch, muß aber in allem Folgenden als Möglichkeit in Betracht gezogen werden.“

Sechs Jahre später schreibt er:

Hausdorff (1914): „Jede Anfangszahl ω_α [jede Kardinalzahl $\aleph_\alpha = \omega_\alpha$] mit Limesindex ist mit ihrem Index α konfinal. Demnach ist sie singulär, falls $\omega_\alpha > \alpha$, und könnte nur regulär sein, wenn $\omega_\alpha = \alpha$, d. h. wenn α eine kritische Zahl für die Normalfunktion ω_α ist [ein Fixpunkt der Alephreihe]. Die erste dieser kritischen Zahlen ist ... der Limes der Zahlen

$$\omega = \omega_0, \omega' = \omega_\omega, \omega'' = \omega_{\omega'}, \dots,$$

... und diese Zahl $\kappa = \omega_\kappa$ ist, obwohl von einer unvorstellbar großen Mächtigkeit, doch noch singulär, da sie als Limes einer [abzählbaren] Folge mit ω konfinal ist. Wenn es also reguläre Anfangszahlen mit Limesindex [schwach unerreichte Kardinalzahlen] gibt (und es ist bisher nicht gelungen, in dieser Annahme einen Widerspruch zu entdecken), so ist die kleinste unter ihnen von einer so exorbitanten Größe, daß sie für die üblichen Zwecke der Mengenlehre kaum jemals in Betracht kommen wird.“

Die „exorbitante Größe“ zeugt von der Irritation, die das Suchen einer schwach unerreichten Kardinalzahl entlang der Alephreihe auslösen kann. Sie liegen außerhalb des Orbits, außerhalb des Gewöhnlichen und des Üblichen. Die Reaktion ist nur natürlich, und es hilft vielleicht dem Leser, der nun unweigerlich mit tiefen mathematischen Existenzfragen konfrontiert werden wird, daß der Entdecker dieser Zahlen auch einige Jahre nach dem Erstkontakt mit dem Phänomen ihnen einen eher abgelegenen Platz in der Mengenlehre prophezeit hat, einfach weil sie einen Platz weit draußen im Universum einnehmen. Für „die üblichen Zwecke“ der Mengenlehre wurden ihre stark unerreichten Verwandten bereits 1930 bedeutsam, als Zermelo „Grenzzahlen“ untersuchte. Sie erscheinen dann als Indikatoren α von Stufen V_α , in denen die ganze übliche Mengenlehre gilt. Weiter erwiesen sie sich als die Kleinsten unter den Großen, und noch viel größere Kardinalzahlen haben erstaunliche Konsequenzen sogar für „die üblichen Zwecke“ nicht nur der Mengenlehre, sondern der Mathematik selber. Sie liefern uns durch eine immer wieder verblüffende intergalaktische Korrespondenz Beweise für Aussagen, die die reellen Zahlen betreffen. Große Kardinalzah-

len können, so der Platoniker, keine Aussagen über kleine Objekte wahr oder falsch machen, aber sie können uns die Möglichkeit in die Hand geben, solche Aussagen zu beweisen. Daß es natürliche Aussagen über $\mathbb{R}$ sind, die mit der Hilfe großer Kardinalzahlen beweisbar werden, ist um so erfreulicher. Wir kommen später noch darauf zurück.

Ist nicht auch ϵ_0 , der erste Fixpunkt der Ordinalzahlexponentiation von einer „exorbitanten Größe“? Die Frage ist immer, welche Operationen man zuläßt, um „groß aus der Sicht von unten“ einen intuitiven Sinn zu geben. Fixpunkte der Ordinalzahlexponentiation und der Alephreihe sind ebenfalls groß in einem gewissen Sinn. Der Unterschied zu unerreichbaren Kardinalzahlen ist, daß wir die Existenz von solchen Fixpunkten mit den üblichen Methoden beweisen können, während wir die Existenz von unerreichbaren Kardinalzahlen mit den üblichen Methoden nicht beweisen können (definitiv nicht, wie man heute weiß). Aber was macht das Übliche der üblichen Methoden aus? Starke Systeme beweisen mehr Aussagen, also ist die platonische Suche nach starken wahren Theorien nur natürlich. Ein Punkt, der große Kardinalzahlen von Aussagen wie (CH) und *non*(CH) unterscheidet – die beide für sich auch Systemverstärkungen bilden würden –, ist, daß große Kardinalzahlen unmittelbar die Ordinalzahlen betreffen, das intrinsisch reichhaltige und unerschöpfliche Grundkonzept der Mengenlehre selber. Warum sollte der unerschöpfliche Pfad der Ordinalzahlen einen Ort nicht erreichen, den er erreichen könnte? Wir hören auch nicht unmittelbar vor ω_1 auf. Der Leser wird sagen: Haben wir nicht bewiesen, daß ω_1 existiert? Natürlich haben wir, aber der Beweis ruht letztendlich auch auf irgendwelchen Annahmen. Die „üblichen Methoden“ generieren in recht subtiler Weise eine schon recht reichhaltige Mengenlehre. Es ist konzeptionell keineswegs klar, daß „es gibt eine schwach unerreichbare Kardinalzahl“ nicht eine ebenso natürliche Annahme ist wie ein mehr oder weniger genau spezifiziertes Bündel von Grundintuitionen über Mengen, das „ ω_1 existiert“ garantiert. Große Kardinalzahlen sind der Kaffee für die Müdigkeit des Üblichen, und sie regen dazu an, über die Üblichkeit des Üblichen genauer nachzudenken.

Die fraglichen Kardinalzahlen heißen heute nicht „exorbitante Kardinalzahlen“ sondern „unerreichbare Kardinalzahlen“, was sehr gut ihr Wesen trifft. Die Bezeichnung stammt von Kuratowski („*nombres cardinaux inaccessibles*“). Der oben kurz erwähnten Idee „ V_α -Stufen bilden Modelle der üblichen Mengenlehre“ folgend gibt man heute Kardinalzahlen mit einer zusätzlichen Eigenschaft den Vorzug:

Definition (*unerreichbare Kardinalzahlen*)

Eine schwach unerreichbare Kardinalzahl κ heißt (*stark*) *unerreichbar*, falls κ ein starker Limes ist, d. h. falls $2^\mu < \kappa$ gilt für alle Kardinalzahlen $\mu < \kappa$.

Zermelo nannte diese Zahlen 1930 *Grenzzahlen*. Im gleichen Jahr untersuchten Sierpiński und Tarski diese Verstärkung des Begriffs einer überabzählbaren regulären Limeskardinalzahl.

Übung

Es gibt unbeschränkt viele Limeskardinalzahlen $\aleph_\lambda$ mit:

- (i) $\aleph_\lambda = \lambda$,
- (ii) $2^\mu < \aleph_\lambda$ für alle Kardinalzahlen $\mu < \aleph_\lambda$,

und diese Kardinalzahlen sind genau die Fixpunkte der Funktion $\beth$, d.h. diejenigen $\aleph_\lambda$ mit $\beth_{\aleph_\lambda} = \aleph_\lambda$.

Die kleinste solche Kardinalzahl hat die Konfinalität ω .

Wir finden also Kardinalzahlen, die viele Bedingungen der starken Unerreichbarkeit erfüllen, „lediglich“ die Regularität macht, wie für die schwache Unerreichbarkeit, Schwierigkeiten. Schwach unerreichbar heißt also „regulär und Fixpunkt der Alephreihe“. Unerreichbar heißt stärker „regulär und Fixpunkt der Bethreihe“.

Ein großes Kardinalzahlaxiom behauptet die Existenz einer Kardinalzahl, die, wenn sie denn existiert, intuitiv sehr groß sein muß. So ist etwa

(U) = „es existiert eine unerreichbare Kardinalzahl“

ein großes Kardinalzahlaxiom, wobei unerreichbar immer stark unerreichbar meint. Was genau eine große Kardinalzahl ist, läßt sich vielleicht gar nicht genau definieren. Jedoch besteht aufgrund der sich ergebenden Struktur der großen Kardinalzahlaxiome heute eine „mathematische Einigkeit“ darüber, welche Aussagen als große Kardinalzahlaxiome zu gelten haben und welche nicht. Der Status aller großen Kardinalzahlaxiome ist dann der folgende. Sei (K) ein großes Kardinalzahlaxiom, etwa (U). Dann gilt:

- (1) Es ist prinzipiell möglich, daß in der üblichen Mengenlehre $\text{non}(K)$ beweisbar ist. Für viele Kardinalzahlaxiome gilt dies aber als sehr unwahrscheinlich. Insbesondere wäre ein Beweis von $\text{non}(U)$ eine Sensation, die mit einem Beweis der Widersprüchlichkeit der üblichen Mengenlehre selbst fast gleichwertig wäre.
- (2) Es ist dagegen prinzipiell unmöglich, (K) innerhalb der üblichen Mengenlehre zu beweisen.
- (3) Es ist sogar prinzipiell unmöglich zu zeigen, daß die Hinzunahme von (K) zur üblichen Mengenlehre keine Widersprüche mit sich bringt.

Der Grund für (2) und (3) ist: (K) beweist die Widerspruchsfreiheit der üblichen Mengenlehre, hat also aufgrund des zweiten Gödelschen Unvollständigkeitssatzes eine über die übliche Mengenlehre hinausgehende logische Kraft. Wir vergleichen die Situation noch einmal mit (CH): Sowohl (CH) als auch $\text{non}(CH)$ können wir zur üblichen Mengenlehre hinzunehmen, ohne die Widerspruchsfreiheit zu zerstören. Und: Wir können beweisen, daß wir die Widerspruchsfreiheit durch die Hinzunahme einer der beiden Aussagen erhalten. Für (K) gilt: *Vermutlich* erhält die Hinzunahme von (K) die Widerspruchsfreiheit. Man kann beweisen, daß dieses *vermutlich* nie in ein *sicherlich* verwandelt werden kann. (Dagegen erhält die Hinzunahme von $\text{non}(K)$ immer die Widerspruchsfreiheit, und daß dies so ist, ist beweisbar.)

Die Rechtfertigung für das Studium der großen Kardinalzahlaxiome ist ihr struktureller Reichtum, die Vielzahl der mit ihnen einhergehenden Phänomene, und ihre Konsequenzen für die Mengenlehre und die Mathematik insgesamt. Die Rechtfertigung für den Glauben an die relative Widerspruchsfreiheit der großen Kardinalzahlaxiome ist gerade dieser Theoriereichtum, und das klare Bild, das man von ihnen heute, nach vielen Jahrzehnten der Untersuchung, zeichnen kann. Die Last der möglichen Widersprüchlichkeit hat man zu tragen gelernt. Schließlich läßt sich ja auch die Widerspruchsfreiheit der üblichen Mengenlehre nicht in der üblichen Mengenlehre beweisen.

Die systemimmanente Unbeweisbarkeit der großen Kardinalzahlaxiome bringt Sprechweisen wie „Glaube an relative Widerspruchsfreiheit“ mit sich, die Mathematiker anderer Gebiete zuweilen an theologische Diskussionen erinnert. Diese Mathematiker glauben aber in einem ganz ähnlichen Sinne auch an die Widerspruchsfreiheit der Arithmetik oder des Mengenlehrefragments, das sie für ihre Beweise brauchen, ohne sich dazu zu bekennen. Sie glauben also an weit mehr, als sie beweisen können.

Zu den strukturellen Besonderheiten der großen Kardinalzahlaxiome gehört ihre Linearität: Von je zwei verschiedenen Kardinalzahlaxiomen (K_1) und (K_2) beweist immer eines die Konsistenz des anderen. So entsteht eine lineare logische Skala. Es zeigt sich, daß natürliche mengentheoretische Aussagen (A) mit dieser Skala gemessen werden können: Für viele (A) hat man bewiesen, daß es genau ein (K) gibt, das die gleiche mengentheoretische Stärke wie (A) hat. Gleiche Stärke heißt genau: Aus der Widerspruchsfreiheit der üblichen Mengenlehre erweitert um (A) folgt die Widerspruchsfreiheit der üblichen Mengenlehre erweitert um (K), und das gleiche gilt mit vertauschten Rollen von (A) und (K). Der zweite Gödelsche Unvollständigkeitssatz macht diese Hierarchie möglich, und wird spätestens jetzt geliebt und nicht mehr als ärgerliche Begrenzung des menschlichen Wissens empfunden. Die großen Kardinalzahlaxiome ordnen die Welt der nicht innerhalb der üblichen Mathematik beweisbaren Aussagen, und man kennt keine Gruppe von Axiomen, die ähnliches zu leisten imstande wäre.

Wir besprechen noch einige weitere große Kardinalzahlen im Umfeld der Unerreichbarkeit. Alle diese Kardinalzahlen gelten heute eher als klein. Das Universum ist ein weites Feld, und die Galaxis, der die Sonne angehört, ist in gewisser Weise bereits riesig groß, aber vom Blickpunkt eines Galaxienhaufens nur ein Staubkorn im All.

Neben (U) gibt es für jede Ordinalzahl α noch die stärkeren Axiome:

(U_α) = „es existieren α -viele unerreichbare Kardinalzahlen“.

Von noch größerer logischer Stärke ist dagegen schon das Axiom:

(U^*) = „es existiert eine unerreichbare Kardinalzahl, die das Supremum von unerreichbaren Kardinalzahlen ist“.

Weiter geht es nach mehreren Zwischenstufen wie etwa (U^*_α) mit:

(U^{**}) = „es existiert eine unerreichbare Kardinalzahl, die das Supremum von unerreichbaren Kardinalzahlen ist, die selbst jeweils das Supremum von unerreichbaren Kardinalzahlen sind“.

Die Iteration derartiger Prinzipien hat Paul Mahlo zwischen 1911 und 1913 in einer Reihe von singulären, ihrer Zeit weit vorauseilenden Arbeiten untersucht (mit Hausdorffs „schwach unerreichbaren Kardinalzahlen“ von 1908 als Ausgangspunkt). Seine Überlegungen führten ihn zu den heute nach ihm benannten Mahlo-Kardinalzahlen, die die erste Stufe „hinter allen“ Operationen des Typs $*$ bilden. Derlei gibt es viele, wie wir sehen werden, auch solche des Typs $*$.

Definition (α -unerreichbare Kardinalzahlen)

Wir definieren κ ist α -unerreichbar für Kardinalzahlen κ und Ordinalzahlen $\alpha \geq 1$ rekursiv wie folgt:

κ heißt 1-*unerreichbar*, falls κ unerreichbar ist.

κ heißt $(\alpha + 1)$ -*unerreichbar*, falls κ unerreichbar und
 $\kappa = \sup \{ \mu < \kappa \mid \mu \text{ ist } \alpha\text{-unerreichbar} \}$.

κ heißt λ -*unerreichbar* für λ Limes, falls κ α -unerreichbar für alle $\alpha < \lambda$ ist.

Die Rekursion dieser Definition können wir streng genommen nicht rechtfertigen, da z. B. $\{ \kappa \mid \kappa \text{ ist unerreichbar} \}$ unbeschränkt in den Ordinalzahlen sein kann, und dann keine Menge mehr ist (vgl. Kapitel 13). Wir können aber dennoch die Rekursion in eine übliche Definition auflösen. Wir definieren also offiziell für $\alpha \geq 1$ und Kardinalzahlen κ :

κ ist α -*unerreichbar*, falls gilt:

κ ist unerreichbar und es gibt eine Folge $\langle X_\gamma \mid 1 \leq \gamma < \alpha \rangle$ mit:

- (i) $X_\gamma \subseteq W(\kappa)$, $\sup(X_\gamma) = \kappa$ für alle $1 \leq \gamma < \alpha$,
- (ii) $X_1 = \{ \mu < \kappa \mid \mu \text{ ist unerreichbar} \}$,
- (iii) $X_{\gamma+1} = \{ \mu \in X_\gamma \mid \sup(X_\gamma \cap W(\mu)) = \mu \}$ für alle $1 \leq \gamma < \alpha$,
- (iv) $X_\lambda = \bigcap_{\gamma < \lambda} X_\gamma$ für alle Limesordinalzahlen $\lambda < \alpha$.

Ist κ α -unerreichbar, so ist κ auch β -unerreichbar für alle $1 \leq \beta < \alpha$.

Der Leser überlegt sich leicht, daß die $(\alpha + 1)$ -unerreichbaren Kardinalzahlen gerade die regulären Fixpunkte der Reihe der α -unerreichbaren Kardinalzahlen sind. In dieser Weise hat Mahlo 1911 die Hierarchie eingeführt.

Der maximale Komplexitätsgrad für ein κ in dieser Hierarchie ist:

„ κ ist κ -unerreichbar“.

[zu maximal: $\langle \mu_\alpha \mid \alpha < \kappa \rangle$ mit $\mu_\alpha =$ „das kleinste α -unerreichbare $\mu < \kappa$ “ ist strikt aufsteigend und konfinal in $W(\kappa)$, sodaß die kleinste $(\kappa + 1)$ -unerreichbare Kardinalzahl sicherlich größer als κ ist.]

Es gibt aber noch trickreiche diagonale Verstärkungen der Hierarchie:

(U^*) „es gibt ein κ -unerreichbares κ , das Limes von μ -unerreichbaren μ ist“,

usw. usw. Mahlo hat nun den Begriff gefunden, der jede derartige Diagonalisierung übersteigt. Er ergibt sich aus heutiger Sicht sehr leicht aus der Durchführung des Themas „ X ist eine große Teilmenge von $W(\kappa)$ “. Dieses Thema ist völlig unabhängig von großen Kardinalzahlen von fundamentaler Bedeutung, und wir wollen ihm in aller Ruhe nachgehen, bevor wir zu Mahlos großartig-großen Kardinalzahlen zurückkehren. Die Motivation ist: Wir suchen einen Begriff, der die Unerreichbarkeits-Hierarchien transzendiert. Wir wollen herausfinden, was „es gibt viele unerreichbare Kardinalzahlen jedes $*$ -Typs unterhalb von κ “ heißen könnte.

Große Teilmengen einer Menge

Sei M eine Menge. Was ist ein guter Begriff für „ X ist eine große Teilmenge von M “? Versuchen wir, einige Anforderungen zu finden, die wir an einen solchen Begriff stellen wollen. Die einfachste ist:

(+) Ist X groß und ist $Y \supseteq X$, so ist auch Y groß.

Darüber sind sich alle einig. Gewagter ist schon:

(++) Sind X und Y groß, so ist auch $X \cap Y$ groß.

Für manche wird sich die duale Form „sind X, Y klein, so ist auch $X \cup Y$ klein“ überzeugender anhören. Die in der endlichen Welt unsinnige Iteration dieser Idee führt zur wahrhaft unentscheidbaren Frage, ab welcher Anzahl eine Menge von Staubkörnern störend wird. Da wir nur mit unendlichen Mengen arbeiten werden, ist die in (++) enthaltene endliche Iteration unproblematisch.

Schwarz-Weiß-Maler oder entscheidungsfreudige Personen fordern:

(+++) Ist X nicht groß, so ist $M - X$ groß, d.h. jedes X ist groß oder klein.

Wie immer ist es die sich erst in der Untersuchung offenbarende Fülle von Konzepten und ihre Interaktion mit anderen Begriffen, die zu guten mathematischen Definitionen führt und sie rechtfertigt. Während sich (++) als unverzichtbar erweist, bleibt die starke Forderung (+++) am besten speziellen Größenbegriffen vorbehalten. Wir definieren demgemäß „ X ist groß in einem gewissen Sinne“ durch „ X ist Element eines Filters“:

Definition (Filter und Ultrafilter)

Sei M eine nichtleere Menge, und sei $F \subseteq \mathcal{P}(M)$.

F heißt ein *Filter auf M* , falls für alle $X, Y \subseteq M$ gilt:

- (i) $\emptyset \notin F, M \in F$,
- (ii) $X \in F$ und $Y \supseteq X$ *folgt* $Y \in F$,
- (iii) $X, Y \in F$ *folgt* $X \cap Y \in F$.

Ein Filter F auf M heißt ein *Ultrafilter auf M* , falls für alle $X \subseteq M$ gilt:

- (iv) $X \in F$ *oder* $M - X \in F$.

Jeder Filter F ist also ein „gewisser Sinn“ von „ X ist groß in einem gewissen Sinne“. Wie so oft haben wir eine Idee extensional gefaßt.

Aus den ersten beiden Eigenschaften folgt: Ist $X \in F$, so ist $M - X \notin F$. Denn andernfalls wäre $X \cap (M - X) = \emptyset \in F$. Die Ultrafilter-Eigenschaft besagt, daß tatsächlich immer entweder X oder $M - X$ Element von F ist.

Anstatt für „groß“ kann man sich auch für „klein“ interessieren. Die zum Begriff eines Filters duale Definition lautet:

Definition (*Ideal und Primideal*)

Sei M eine nichtleere Menge, und sei $I \subseteq \mathcal{P}(M)$.

I heißt ein *Ideal auf M* , falls für alle $X, Y \subseteq M$ gilt:

- (i) $\emptyset \in I$, $M \notin I$,
- (ii) $X \in I$ und $Y \subseteq X$ folgt $Y \in I$,
- (iii) $X, Y \in I$ folgt $X \cup Y \in I$.

Ein Ideal I auf M heißt ein *Primideal auf M* , falls für alle $X \subseteq M$ gilt:

- (iv) $X \in I$ oder $M - X \in I$.

Es ist Geschmackssache, ob man Filtern oder Idealen den Vorzug gibt. Die erzielten Resultate übertragen sich immer auf den vernachlässigten der beiden Begriffe. Im folgenden bevorzugen wir zumeist „groß“ statt „klein“.

Beispiele für Filter sind: Sei M eine unendliche Menge. Dann ist

$$F_{\text{Fin}} = \{ X \subseteq M \mid M - X \text{ ist endlich} \}$$

ein Filter auf M .

Ist $|M| = \kappa \geq \omega$, und ist $\mu \leq \kappa$ eine unendliche Kardinalzahl, so ist

$$F_{<\mu} = \{ X \subseteq M \mid |M - X| < \mu \}$$

ein Filter auf M .

Das zweite Beispiel spezialisiert sich für $\mu = \omega$ zum ersten.

Ist M eine Menge und $f: M \rightarrow W(\kappa)$ bijektiv, so übersetzt f einen Filter F auf M in einen Filter F' auf $W(\kappa)$: $F' = \{ X \subseteq W(\kappa) \mid f^{-1}X \in F \}$. Die zugrundeliegenden Mengen M werden daher zumeist Kardinalzahlen sein. Ein weiterer Grund, große Teilmengen von $W(\kappa)$ zu untersuchen, ist die ihnen innewohnende Ordnung. Ist κ eine Kardinalzahl, so ist:

$$F_E = \{ X \subseteq W(\kappa) \mid X \supseteq W(\kappa) - W(\alpha) \text{ für ein } \alpha < \kappa \}$$

ein Filter auf M . Groß wird hier als „ X enthält ein Endstück von $W(\kappa)$ “ interpretiert. Ist $\kappa \geq \omega$ regulär, so ist der Endstück-Filter identisch mit $F_{<\kappa}$. Für singuläre $\kappa \geq \omega$ ist $F_E \subset F_{<\kappa}$.

club-Mengen

„ X ist unbeschränkt in $W(\kappa)$ “ verletzt die Schnitteigenschaft von Filtern. Eine Verstärkung der Unbeschränktheit führt zu dem vielleicht interessantesten Filter auf $W(\kappa)$ – und zu den Mahlo-Kardinalzahlen.

Definition (*unbeschränkt, abgeschlossen, club*)

Sei λ eine Limesordinalzahl, und sei $X \subseteq W(\lambda)$.

- (i) X heißt *unbeschränkt in λ* , falls $\lambda = \sup(X)$.
- (ii) X heißt *abgeschlossen in λ* , falls für alle Limesordinalzahlen $\alpha < \lambda$ gilt:
 $\sup(X \cap W(\alpha)) = \alpha$ folgt $\alpha \in X$.
- (iii) X heißt *club in λ* , falls X abgeschlossen und unbeschränkt in λ ist.

Der Ausdruck „club“ ist eine Verballhornung von engl. closed-unbounded, und auch im Deutschen bequem verwendbar. Wir konzentrieren uns im folgenden auf reguläre Kardinalzahlen $\lambda \geq \omega_1$. Vieles gilt allgemeiner auch für Limesordinalzahlen λ mit $\text{cf}(\lambda) \geq \omega_1$.

Übung

Sei $\kappa \geq \omega_1$ regulär, und sei $C \subseteq W(\kappa)$. Dann sind äquivalent:

- (i) C ist club in κ .
- (ii) Es gibt eine strikt aufsteigende stetige Folge $f : W(\kappa) \rightarrow W(\kappa)$ mit $C = \text{rng}(f)$.

Die folgenden auf den ersten Blick verblüffenden Fakten werden oft verwendet. Beliebige Funktionen auf regulären überabzählbaren Kardinalzahlen holen sich club-of selbst ein, und stetige kofinale Funktionen haben club-viele Fixpunkte:

Übung

Sei $\kappa \geq \omega_1$ regulär, und sei $f : W(\kappa) \rightarrow W(\kappa)$ eine Funktion. Dann gilt:

- (i) $\{ \alpha < \kappa \mid f''\alpha \subseteq W(\alpha) \}$ ist club in κ .
[Abgeschlossen ist klar. Zu unbeschränkt betrachte $\alpha_\omega = \sup_{n < \omega} \alpha_n$ für $\alpha_{n+1} = \sup f''(\alpha_n + 1)$ für $n < \omega$, mit $\alpha_0 < \kappa$ beliebig.]
- (ii) Ist f stetig und kofinal in κ , so ist $\{ \alpha < \kappa \mid f(\alpha) = \alpha \}$ club in κ .

Die monotone Aufzählung $\langle \gamma_\alpha \mid \alpha < \kappa \rangle$ einer club-Menge $C \subseteq W(\kappa)$ ist stetig und kofinal in κ und hat also club-viele Fixpunkte. Club-Mengen sind sehr stabil gegenüber derartigen Ausdünnungsprozessen.

Jede in κ unbeschränkte Menge $X \subseteq W(\kappa)$ läßt sich kanonisch zu einer club-Menge erweitern, indem man alle von X angesteuerten Punkte zu X hinzunimmt:

Definition (Limespunkt und Abschluß)

Sei X eine Menge von Ordinalzahlen, und sei λ eine Limesordinalzahl.

λ heißt (*innerer*) *Limespunkt* von X , falls $\lambda < \sup(X)$ und $\lambda = \sup(X \cap W(\lambda))$.

Wir setzen:

$\text{Lim}(X) = \{ \lambda \mid \lambda \text{ ist innerer Limespunkt von } X \}$.

$X \cup \text{Lim}(X)$ heißt der *Abschluß* von X .

Offenbar ist ein unbeschränktes $X \subseteq W(\kappa)$ genau dann club in κ , wenn $\text{Lim}(X) \subseteq X$ gilt. Die Mengen $\text{Lim}(X)$ und $X \cup \text{Lim}(X)$ sind club in κ für alle unbeschränkten $X \subseteq W(\kappa)$.

Wir zeigen nun durch ein hübsches Verflechtungsargument, daß die club-Mengen von regulären $\kappa \geq \omega_1$ einen Filter bilden.

Satz (Durchschnitt von club-Mengen)

Sei $\kappa \geq \omega_1$ regulär, und seien C, D club in κ .

Dann ist $C \cap D$ club in κ .

Beweis

$C \cap D$ ist abgeschlossen in κ :

Sei $\lambda = \sup(C \cap D \cap W(\lambda))$ für eine Limesordinalzahl $\lambda < \kappa$.

Dann ist $\lambda = \sup(C \cap W(\lambda))$, also $\lambda \in C$ wegen C abgeschlossen.

Analog gilt $\lambda \in D$. Also $\lambda \in C \cap D$.

$C \cap D$ ist unbeschränkt in κ :

Wir definieren durch Rekursion über $\alpha < \kappa$:

$\gamma_{2\alpha} =$ „das kleinste $\gamma \in C$ mit $\gamma > \gamma_\beta$ für alle $\beta < 2\alpha$ “,

$\gamma_{2\alpha+1} =$ „das kleinste $\gamma \in D$ mit $\gamma > \gamma_\beta$ für alle $\beta \leq 2\alpha$ “.

Wegen κ regulär ist $\gamma_\alpha < \kappa$ für alle $\alpha < \kappa$.

Sei $E = \{ \gamma_\lambda \mid \lambda < \kappa, \lambda \text{ Limes} \}$.

Dann ist $E \subseteq C \cap D$ club in $W(\kappa)$, wovon der Leser sich mit Freuden überzeugen kann. Insbesondere ist $C \cap D$ also unbeschränkt in κ .

Wir beweisen hier (und in den folgenden Argumenten) viel mehr als wir brauchen. Der Satz bleibt richtig mit der Voraussetzung $\text{cf}(\kappa) \geq \omega_1$ (!) [abgeschlossen ist klar, und eine ω -Rekursion genügt für die Unbeschränktheit]. Wir begnügen uns hier aber mit der Betrachtung von regulären $\kappa \geq \omega_1$, und konstruieren ganze club-Mengen in Schnitten, um die Unbeschränktheit zu zeigen.

Allgemeiner gilt:

Übung

Sei $\kappa \geq \omega_1$ regulär, und sei $\beta < \kappa$. Weiter seien C_α club in κ für $\alpha < \beta$.

Dann ist $\bigcap_{\alpha < \beta} C_\alpha$ club in κ .

[Definiere rekursiv für $\delta < \kappa$, $\alpha < \beta$:

$\gamma_{\beta \cdot \delta + \alpha} =$ „das kleinste $\gamma \in C_\alpha$ mit $\gamma > \gamma_\eta$ für alle $\eta < \beta \cdot \delta + \alpha$ “

Sei $E = \{ \gamma_{\beta \cdot \lambda} \mid \lambda < \kappa, \lambda \text{ Limes} \}$. Dann ist E club und $E \subseteq C_\alpha$ für $\alpha < \beta$.]

Definition (μ -vollständige Filter und Ideale)

Sei F ein Filter und sei $\mu \geq \omega_1$ eine Kardinalzahl.

F heißt μ -vollständig, falls $\bigcap X \in F$ für alle $X \subseteq F$ mit $|X| < \mu$.

Analog heißt ein Ideal I μ -vollständig, falls $\bigcup X \in I$ für alle $X \subseteq I$ mit $|X| < \mu$.

μ -Vollständigkeit ist also eine Verstärkung der Schnitteigenschaft von Filtern. Speziell meint ω_1 -vollständig, daß Filter stabil sind unter abzählbaren Schnittbildungen.

Definition (club-Filter, dünnes oder nichtstationäres Ideal)

Sei $\kappa \geq \omega_1$ regulär. Wir setzen:

$\mathcal{C}_\kappa = \{ X \subseteq W(\kappa) \mid \text{es gibt ein } C \subseteq W(\kappa) \text{ mit } C \text{ club in } \kappa \text{ und } C \subseteq X \}$,

$\mathcal{I}_\kappa = \{ X \subseteq W(\kappa) \mid W(\kappa) - X \in \mathcal{C}_\kappa \}$.

$\mathcal{C}_\kappa$ heißt der club-Filter auf κ ,

$\mathcal{I}_\kappa$ heißt das dünne oder nichtstationäre Ideal auf κ .

Eine Teilmenge X von $W(\kappa)$ heißt konsequenterweise *dünn* oder *nichtstationär* in κ , falls $X \in \mathcal{I}_\kappa$. Die Bezeichnung „nichtstationär“ wird gleich klar werden.

Wir haben mit der obigen Übung gezeigt:

Korollar (κ -Vollständigkeit des club-Filters)

Sei $\kappa \geq \omega_1$ regulär.

Dann ist $\mathcal{C}_\kappa$ ein κ -vollständiger Filter auf κ .

Dual gilt: $\mathcal{I}_\kappa$ ist ein κ -vollständiges Ideal auf κ .

Es gilt nun eine noch stärkere Durchschnittseigenschaft mit verblüffenden Konsequenzen. Hierzu definieren wir:

Definition (*diagonaler Durchschnitt und diagonale Vereinigung*)

Sei $\kappa \geq \omega$ eine Kardinalzahl, und sei $X_\alpha \subseteq W(\kappa)$ für $\alpha < \kappa$.

Dann sind der *diagonale Durchschnitt* $\Delta_{\alpha < \kappa} X_\alpha$ und die *diagonale Vereinigung* $\nabla_{\alpha < \kappa} X_\alpha$ der Folge $\langle X_\alpha \mid \alpha < \kappa \rangle$ definiert durch:

$$\Delta_{\alpha < \kappa} X_\alpha = \{ \alpha < \kappa \mid \alpha \in X_\beta \text{ für alle } \beta < \alpha \},$$

$$\nabla_{\alpha < \kappa} X_\alpha = \{ \alpha < \kappa \mid \alpha \in X_\beta \text{ für ein } \beta < \alpha \}.$$

Der diagonale Schnitt und die diagonale Vereinigung haben eine einfache geometrische Deutung. Für $X \subseteq W(\kappa)$ seien

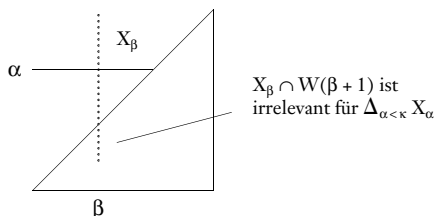
$$X^{+, \alpha} = X \cup W(\alpha + 1), \text{ und}$$

$$X^{-, \alpha} = X - W(\alpha + 1).$$

Dann gilt:

$$\Delta_{\alpha < \kappa} X_\alpha = \bigcap_{\alpha < \kappa} X_\alpha^{+, \alpha},$$

$$\nabla_{\alpha < \kappa} X_\alpha = \bigcup_{\alpha < \kappa} X_\alpha^{-, \alpha}.$$



Anders: Wir interessieren uns bei den diagonalen Operationen nur für das „linke obere Dreieck“ $D = \{ (\beta, \alpha) \mid \beta < \alpha < \kappa \}$ von $W(\kappa) \times W(\kappa)$, ausschließlich der Diagonalen, und bilden Schnitt und Vereinigung bzgl. D . Ein $\alpha < \kappa$ wird in den diagonalen Schnitt genau dann aufgenommen, wenn es zu allen X_β mit $\beta < \alpha$ gehört. Analoges gilt mit „einem“ statt „allen“ für die diagonale Vereinigung.

Der diagonale Schnitt taucht zum ersten Mal auf in [Veblen 1908].

Ganz anders als der Schnitt und die Vereinigung hängen die diagonalen Versionen von der Aufzählung der zu schneidenden Mengen ab. Dennoch sind sie „modulo club“ eindeutig bestimmt:

Übung

Sei $\kappa \geq \omega_1$ regulär, und seien $X_\alpha, Y_\alpha \subseteq W(\kappa)$ für $\alpha < \kappa$ mit

$\{ X_\alpha \mid \alpha < \kappa \} = \{ Y_\alpha \mid \alpha < \kappa \}$. Dann sind die symmetrischen Differenzen

$\Delta_{\alpha < \kappa} X_\alpha \Delta \Delta_{\alpha < \kappa} Y_\alpha$ und $\nabla_{\alpha < \kappa} X_\alpha \Delta \nabla_{\alpha < \kappa} Y_\alpha$ dünn in κ .

[Die Menge $C = \{ \alpha < \kappa \mid \{ X_\beta \mid \beta < \alpha \} = \{ Y_\beta \mid \beta < \alpha \} \}$ ist club in κ .]

Club-Folgen kann nun auch ein diagonaler Schnitt nicht beeindrucken. Das Ergebnis eines solchen Schnitts ist wieder eine club-Menge:

Satz (*über diagonale Schnitte*)

Sei $\kappa > \omega$ regulär, und seien C_α club in κ für $\alpha < \kappa$.

Dann ist $\Delta_{\alpha < \kappa} C_\alpha$ club in κ .

Beweis

Sei $D = \Delta_{\alpha < \kappa} C_\alpha$.

D ist abgeschlossen in κ :

Sei $\lambda = \sup(D \cap W(\lambda))$ für eine Limesordinalzahl $\lambda < \kappa$.

Es gilt $D \cap W(\lambda) - W(\beta + 1) \subseteq C_\beta$ für alle $\beta < \lambda$.

Also ist λ Limespunkt von C_β für $\beta < \lambda$.

Dann ist aber $\lambda \in C_\beta$ für alle $\beta < \lambda$, also $\lambda \in D$.

D ist unbeschränkt in κ :

Wir setzen $\gamma_0 = 0$ und definieren durch Rekursion über $1 \leq \alpha < \kappa$:

$\gamma_\alpha =$ „das kleinste $\gamma \in \bigcap_{\beta < \alpha} C_\beta$ mit $\gamma > \gamma_\beta$ für alle $\beta < \alpha$ “.

Wegen κ regulär ist $\gamma_\alpha < \kappa$ für alle $\alpha < \kappa$.

Weiter ist $\langle \gamma_\alpha \mid \alpha < \kappa \rangle$ strikt aufsteigend und stetig.

Also ist $E = \{ \alpha < \kappa \mid \gamma_\alpha = \alpha \}$ club in κ .

— Nach Konstruktion gilt $E \subseteq D$, also ist D unbeschränkt in κ .

Man kann die Rekursion auch so definieren:

$\gamma_\alpha =$ „das kleinste γ mit: für alle $\beta < \alpha$ existiert ein $\delta \in C_\beta$, $\delta \leq \gamma$ mit $\delta > \gamma_{\beta'}$ für alle $\beta' < \alpha$.“

Dann ist wieder $E = \{ \alpha < \kappa \mid \gamma_\alpha = \alpha \}$ club in κ und $E \subseteq D$. Der Beweis ruht dann nicht mehr auf dem vorab bewiesenen Resultat, daß der Schnitt von weniger als κ -vielen club-Mengen wieder club ist, und die Schnitt-Sätze werden zu einfachen Korollaren zum Satz über diagonale Durchschnitte. Eine schrittweise Näherung an die maximale Abgeschlossenheitseigenschaft von club-Mengen erscheint aber natürlicher.

Definition (*normaler Filter und normales Ideal*)

Sei $\kappa \geq \omega$ eine Kardinalzahl und sei F ein Filter auf $W(\kappa)$.

F heißt *normal*, falls F abgeschlossen unter diagonalen Durchschnitten ist,

d. h. es gilt $\Delta_{\alpha < \kappa} X_\alpha \in F$ für alle $\langle X_\alpha \mid \alpha < \kappa \rangle$ mit $X_\alpha \in F$ für alle $\alpha < \kappa$.

Analog heißt ein Ideal I auf $W(\kappa)$ *normal*, falls I abgeschlossen unter diagonalen Vereinigungen ist.

Übung

Sei $\kappa \geq \omega$ eine Kardinalzahl und sei F ein normaler Filter auf $W(\kappa)$.

Weiter gelte $W(\kappa) - W(\alpha) \in F$ für alle $\alpha < \kappa$.

Dann ist F κ -vollständig, und κ ist regulär und überabzählbar.

Wir haben gezeigt:

Korollar (*Normalität des club-Filters*)

Sei $\kappa \geq \omega_1$ regulär. Dann ist $\mathcal{C}_\kappa$ ein normaler Filter auf $W(\kappa)$.

Dual gilt: $\mathcal{I}_\kappa$ ist ein normales Ideal für reguläre $\kappa \geq \omega_1$.

Stationäre Mengen

Eine natürliche Frage für jedes reguläre $\kappa \geq \omega_1$ ist: Ist jede Menge $X \subseteq W(\kappa)$ entweder in $\mathcal{C}_\kappa$ oder in $\mathcal{I}_\kappa$? Mit anderen Worten: Ist $\mathcal{C}_\kappa$ ein Ultrafilter? Für den Fall $\kappa = \omega_2$ ist die Frage leicht zu verneinen:

Übung

Sei $S_0 = \{ \alpha < \omega_2 \mid \text{cf}(\alpha) = \omega \}$ und $S_1 = \{ \alpha < \omega_2 \mid \text{cf}(\alpha) = \omega_1 \}$.

Zeigen Sie $S_0, S_1 \notin \mathcal{C}_{\omega_2} \cup \mathcal{I}_{\omega_2}$.

[Ist C club in ω_2 , so ist $C \cap S_0 \neq \emptyset$ und $C \cap S_1 \neq \emptyset$.]

Allgemein gilt: Sind $\omega \leq \mu < \kappa$ regulär, so ist $S_\mu = \{ \alpha < \kappa \mid \text{cf}(\alpha) = \mu \} \notin \mathcal{C}_\kappa \cup \mathcal{I}_\kappa$.

Definition (*stationär*)

Sei $\kappa \geq \omega_1$ regulär, und sei $S \subseteq W(\kappa)$. S heißt:

- (i) *stationär*, falls $S \notin \mathcal{I}_\kappa$,
- (ii) *kostationär*, falls $W(\kappa) - S$ stationär ist,
- (iii) *stationär-kostationär*, falls S stationär und kostationär ist.

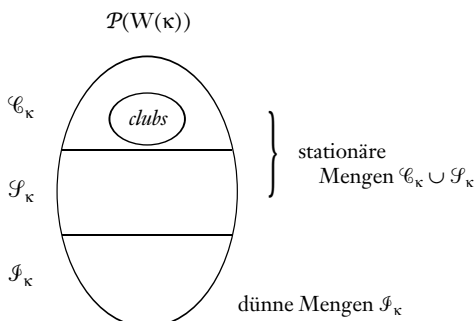
Wir setzen:

$$\mathcal{S}_\kappa = \{ X \subseteq W(\kappa) \mid X \text{ ist stationär-kostationär} \} = \mathcal{P}(W(\kappa)) - (\mathcal{C}_\kappa \cup \mathcal{I}_\kappa).$$

Die stationär-kostationären Mengen sind gerade die Mengen, die weder club-groß noch club-klein sind. Die Mengen $S_0, S_1 \subseteq W(\omega_2)$ aus der obigen Übung sind Beispiele für Mengen in $\mathcal{S}_{\omega_2}$. Sie sind nicht „club-meßbar“.

Ganz $\mathcal{P}(W(\kappa))$ zerfällt, wie einst Gallien bei Cäsar, in drei Teile, nämlich in die paarweise disjunkten Systeme $\mathcal{C}_\kappa, \mathcal{I}_\kappa, \mathcal{S}_\kappa$.

Den Begriff der stationären Menge hat Mahlo untersucht (1911), die Bezeichnung „stationär“ stammt von Gérard Bloch (1953). Die Wortwahl ist motiviert durch die folgenden Fakten, die ständig gebraucht werden:



Übung

Sei $\kappa \geq \omega_1$ regulär, und sei $S \subseteq W(\kappa)$. Dann sind äquivalent:

- (i) S ist stationär in κ .
- (ii) $S \cap C \neq \emptyset$ für jede club-Menge C in κ .

Übung

Seien $\kappa \geq \omega_1$ regulär, $C \in \mathcal{C}_\kappa$ und S stationär in κ .
Dann ist $C \cap S$ stationär in κ .

Wir werden unten zeigen, daß $\mathcal{S}_\kappa$ stets sehr reich ist. Bislang haben wir noch nicht ausgeschlossen, daß $\mathcal{S}_{\omega_1} = \emptyset$. Zuvor besprechen wir noch eine erstaunliche Konsequenz der Abgeschlossenheit von club-Mengen unter diagonalen Schnitten: Regressive Funktionen, die auf beträchtlichen Mengen definiert sind, sind auf beträchtlichen Mengen konstant.

Definition (*regressive Funktionen*)

Seien $\kappa \geq \omega_1$ regulär, $S \subseteq W(\kappa)$ und $f : S \rightarrow W(\kappa)$.
 f heißt *regressiv auf S* , falls $f(\alpha) < \alpha$ für alle $\alpha \in S - \{0\}$ gilt.

Korollar (*Satz von Fodor*)

Sei $\kappa > \omega$ regulär, und sei $S \subseteq W(\kappa)$ stationär.
Weiter sei $f : S \rightarrow W(\kappa)$ regressiv auf S .
Dann existiert ein $\sigma < \kappa$ derart, daß $\{\alpha \in S \mid f(\alpha) = \sigma\}$ stationär in κ ist.

Dieser Satz erscheint in [Fodor 1956], Vorstufen davon finden sich in [Alexandrov / Urysohn 1929].

Beweis

O. E. gilt $0 \notin S$. *Annahme*, ein $\sigma < \kappa$ wie im Satz existiert nicht.

Dann ist $A_\sigma = f^{-1} \{\sigma\}$ dünn für alle $\sigma < \kappa$ und damit ist $\nabla_{\sigma < \kappa} A_\sigma$ dünn.

Wegen f regressiv ist aber $A_\sigma \subseteq W(\kappa) - W(\sigma + 1)$ für alle $\sigma < \kappa$. Also ist

– $S = \bigcup_{\sigma < \kappa} A_\sigma = \nabla_{\sigma < \kappa} A_\sigma$ dünn, *im Widerspruch* zu S stationär.

Es gibt also unterhalb der Diagonalen keine streng monoton wachsenden Funktionen auf einer regulären Kardinalzahl $\kappa \geq \omega_1$! Man vergleiche dies mit der Funktion $f : \mathbb{N} - \{0\} \rightarrow \mathbb{N}$ mit $f(n) = n - 1$.

Der Satz von Fodor ist im Grunde nur eine Umformulierung der Abgeschlossenheit von club-Mengen unter diagonalen Schnitten:

Übung

Sei $\kappa \geq \omega_1$ regulär, und seien $C_\alpha \subseteq W(\kappa)$ club in κ für $\alpha < \kappa$.

Zeigen Sie mit Hilfe des Satzes von Fodor, daß $\Delta_{\alpha < \kappa} C_\alpha$ club in κ ist.

[*Andernfalls* ist f mit $f(\alpha) = \sup(\{\beta < \alpha \mid \beta \in C_\alpha\})$ regressiv auf einer stationären Menge.]

Ulam-Matrizen

Bevor wir nun zu den Mahlo-Kardinalzahlen kommen, betrachten wir noch die bislang etwas dunkel gebliebenen stationär-kostationären Mengen. Speziell für $\kappa = \omega_1$ ist die Existenz von derartigen Mengen ein nichttriviales, nicht-elementares Resultat. Allgemein existiert eine Zerlegung einer stationären Menge auf κ in κ -viele paarweise disjunkte stationäre Teile.

Von den definierbaren Konfinalitäts-Aussonderungen $S_\mu = \{ \alpha < \kappa \mid \text{cf}(\alpha) = \mu \}$ für reguläre $\omega \leq \mu < \kappa$ und reguläre $\kappa \geq \omega_2$ einmal abgesehen, müssen sehr abstrakte Argumente verwendet werden, um überhaupt stationär-kostationäre, also halb-große-halbkleine Mengen zu finden. Die folgende Konstruktion gehört zu den Juwelen der kombinatorischen Mengenlehre. Sie stammt von Stanisław Ulam und findet sich in einem etwas anderen Zusammenhang in einer epochalen Arbeit von Ulam aus dem Jahr 1930.

Definition (*Ulam-Matrix*)

Sei $\mu \geq \omega$ eine Kardinalzahl, und sei $\kappa = \mu^+$.

$\langle X_{v,\alpha} \mid v < \mu, \alpha < \kappa \rangle$ heißt eine *Ulam-Matrix* auf κ , falls gilt:

- (i) $X_{v,\alpha} \subseteq W(\kappa)$ für alle $v < \mu, \alpha < \kappa$,
- (ii) $\bigcup_{v < \mu} X_{v,\alpha} = W(\kappa) - W(\alpha)$ für alle $\alpha < \kappa$,
- (iii) $X_{v,\alpha} \cap X_{v,\beta} = \emptyset$ für alle $v < \mu$ und alle $\alpha, \beta < \kappa$ mit $\alpha \neq \beta$.

Die Spalten der Höhe μ einer Ulam-Matrix ergänzen sich also zu fast ganz $W(\kappa)$, während die κ -vielen Einträge auf jeder Zeile der Matrix paarweise disjunkt sind.

Satz (*Existenz von Ulam-Matrizen für Nachfolgerkardinalzahlen*)

Sei $\mu \geq \omega$ eine Kardinalzahl, und sei $\kappa = \mu^+$.

Dann existiert eine Ulam-Matrix auf κ .

Beweis

Für $\beta < \kappa$ sei

$g_\beta = \text{„ein surjektives } g : W(\mu) \rightarrow W(\beta + 1)\text{“}.$

Wir setzen für $v < \mu$ und $\alpha < \kappa$:

$X_{v,\alpha} = \{ \beta < \kappa \mid g_\beta(v) = \alpha \}.$

– Dann ist $\langle X_{v,\alpha} \mid v < \mu, \alpha < \kappa \rangle$ wie gewünscht.

Das Argument ist ebenso originell wie irritierend einfach. Die Konsequenzen für die Maßtheorie sind weitreichend (s. u.). Für hier genügt uns die Anwendung einer Ulam-Matrix auf stationäre Mengen: In jeder Spalte einer Ulam-Matrix auf $\kappa = \mu^+$ sitzt zumindest eine stationäre Menge, und ein einfaches Stabilitätsargument für den zugehörigen Zeilenindex liefert dann den folgenden Zerlegungssatz.

Korollar (*Existenz von stationären Zerlegungen*)

Sei $\mu \geq \omega$ eine Kardinalzahl, und sei $\kappa = \mu^+$.

Dann existiert eine Menge $\mathfrak{S}$ von paarweise disjunkten in κ stationären Mengen mit $|\mathfrak{S}| = \kappa$.

Beweis

Sei $\langle X_{v,\alpha} \mid v < \mu, \alpha < \kappa \rangle$ eine Ulam-Matrix auf κ .

Für $\alpha < \kappa$ sei:

$f(\alpha) =$ „das kleinste $v < \mu$ mit $X_{v,\alpha}$ stationär“.

Ein solches v existiert, denn $\bigcup_{v < \mu} X_{v,\alpha} = W(\kappa) - W(\alpha)$ ist nicht dünn für alle $\alpha < \kappa$, also nicht die Vereinigung von μ -vielen dünnen Teilmengen von $W(\kappa)$.

Wegen $\text{cf}(\kappa) = \kappa > \mu$ existiert ein $A \subseteq W(\kappa)$ mit $|A| = \kappa$ und ein $v^* < \mu$ mit $f(\alpha) = v^*$ für alle $\alpha \in A$.

- Dann ist $\mathfrak{S} = \{ X_{v^*,\alpha} \mid \alpha \in A \}$ wie gewünscht.

Eine Variation der Konstruktion zeigt:

Übung

Sei $\kappa = \mu^+$ für eine Kardinalzahl $\mu \geq \omega$, und sei $S \subseteq W(\kappa)$ stationär in κ .

Dann existiert eine Menge $\mathfrak{S}$ von paarweise disjunkten Mengen mit:

- (i) $T \subseteq S$ für alle $T \in \mathfrak{S}$,
- (ii) T ist stationär in κ für alle $T \in \mathfrak{S}$,
- (iii) $|\mathfrak{S}| = \kappa$.

Ein analoges Zerlegungsergebnis wie im Satz oben und in der Übung ist nicht nur für Nachfolgerkardinalzahlen $\kappa = \mu^+$, sondern auch für reguläre Limeskardinalzahlen κ gültig, wie Solovay 1971 gezeigt hat.

Speziell ist also $\mathcal{S}_{\omega_1}$ sehr reichhaltig, und es ergeben sich eine Menge von schwierigen Fragen über die Struktur des club-Filters auf ω_1 , oder allgemeiner auf regulären Kardinalzahlen κ . Die Energie, die in den letzten Jahrzehnten in der mengentheoretischen Forschung auf die Untersuchung von $\mathcal{C}_{\omega_1}$ verwendet wurde, dürfte in etwa der Energie entsprechen, die in die Kontinuumshypothese investiert wurde. Zwei Strukturfragen, die man über die stationär-kstationären Mengen $\mathcal{S}_{\omega_1}$ auf ω_1 stellen kann, sind:

- (1) Existiert eine Folge $\langle S_\alpha \mid \alpha < \omega_2 \rangle$ in $\mathcal{S}_{\omega_1}$ mit:
 $S_\alpha \cap S_\beta$ ist dünn für alle $\alpha, \beta < \omega_2$ mit $\alpha \neq \beta$?
- (2) Existiert eine Folge $\langle S_\alpha \mid \alpha < \omega_1 \rangle$ in $\mathcal{S}_{\omega_1}$ mit:
Für alle $X \in \mathcal{S}_{\omega_1}$ existiert ein $\alpha < \omega_1$ mit $X - S_\alpha$ ist dünn?

Ist die Antwort auf die erste Frage *nein*, so sagt man, daß $\mathcal{C}_{\omega_1}$ ω_2 -saturiert ist. Ist die Antwort auf die zweite Frage *ja*, so sagt man, daß $\mathcal{C}_{\omega_1}$ ω_1 -dicht ist.

Übung

Ist $\mathcal{C}_{\omega_1}$ ω_1 -dicht, so ist $\mathcal{C}_{\omega_1}$ ω_2 -saturiert.

Die beiden Fragen lassen sich in der üblichen Mengenlehre nicht beantworten, wobei der notorische Zusatz „es sei denn, die Mengenlehre ist widerspruchsvoll“ hier heißen muß: „es sei denn, die um große Kardinalzahlaxiome erweiterte Mengenlehre ist widerspruchsvoll“. Im Gegensatz zu (CH) und *non*(CH) – und analog zu *non*(SCH) – erfordert die Konstruktion von Modellen, in denen $\mathcal{C}_{\omega_1}$ ω_2 -saturiert oder sogar ω_1 -dicht ist, große Kardinalzahlen als Baumaterial. Umgekehrt liefern die beiden kombinatorischen Prinzipien große Kardinalzahlen in sog. inneren Modellen, die ähnlich wie Gödels Modell L für (CH) sorgfältig von unten aufgebaut sind, und dabei aber im Gegensatz zu L von V Objekte stehlen, die großen Kardinalzahlen zugeordnet sind. Die verwendeten und zurückgewonnenen Kardinalzahlen sind viel, viel größer als unerreichbare und meßbare Kardinalzahlen, es handelt sich um die sog. Woodin-Kardinalzahlen (benannt nach dem Mengentheoretiker Hugh Woodin).

Die Saturiertheits- und Dichtheits-Resultate über den club-Filter auf ω_1 stammen aus dem Ende des 20. Jahrhunderts. Sie erfordern sehr aufwendige Techniken, und eine Vielzahl von Forschern hat zu dem Gebäude beigetragen, das derartige Untersuchungen ermöglicht. Wir verweisen den Leser hier auf den Kommentar zum Literaturverzeichnis, und von dort aus auf die neuere Lehrbuchliteratur, die sich solchen Spezialfragen widmet.

Mahlo-Kardinalzahlen

Nach diesem kombinatorischen Interregnum können wir nun Mahlo-Kardinalzahlen definieren, ohne daß diese Definition einfach vom Himmel fallen würde. Wir fordern nun nicht nur wie etwa für „ κ ist 2-unerreichbar“, daß unbeschränkt viele unerreichbare Kardinalzahlen unterhalb von κ ihr platonisches Dasein fristen, sondern daß unerreichbare Kardinalzahlen unterhalb von κ „beträchtlich“ – oft vorkommen. Club-oft können wir sicher nicht fordern, weil jede club-Menge ein Element der Konfinalität ω enthält (!), während unerreichbare Kardinalzahlen regulär sind. Das Beste, was man haben kann, ist stationär-oft. Und das tuts:

Definition (*Mahlo-Kardinalzahl*)

Sei $\kappa \geq \omega_1$ regulär. κ heißt eine *Mahlo-Kardinalzahl*, falls gilt:

$\{ \mu < \kappa \mid \mu \text{ ist unerreichbar} \}$ ist stationär in κ .

Paul Mahlo (1911): „Wir könnten unser Untersuchungsgebiet [Hierarchien von unerreichbaren Kardinalzahlen] noch beliebig erweitern und kämen doch immer zum gleichen Resultate, so lange wir nur mit den sich sukzessive ergebenden Zahlen operieren. Aber zugleich kann man die Gesamtheit der so zu erzeugenden Zahlen in eine Klasse zusammenfassen und ihrer Reihe eine Zahl $\rho_{1,0}$ [die kleinste Mahlo-Kardinalzahl] mit folgender Eigenschaft ... folgen lassen: wie auch eine Reihe von Zahlen $\alpha_0, \alpha_1, \dots$ mit $\rho_{1,0}$ als Limes aufgestellt wird, stets sollen $\rho_{1,0}$ verschiedene Abschnitte davon mit π_0 -Zahlen [regulären Zahlen] als Limites existieren... ebenso können wir weitere Zahlen von entsprechender Eigenschaft wie $\rho_{1,0}$ definieren, die wir ρ_0 -Zahlen [Mahlo-Kardinalzahlen] nennen.“

κ ist also nach Mahlos Definition eine ρ_0 -Zahl, falls gilt: für alle in κ unbeschränkten $A \subseteq W(\kappa)$ ist $|\{\mu < \kappa \mid \mu \text{ ist regulär und ein Limespunkt von } A\}| = \kappa$. Diese Definition ist gleichwertig mit „ $\{\mu < \kappa \mid \mu \text{ ist schwach unerreichbar}\}$ ist stationär in κ “.

Der Ausdruck „lassen wir folgen“ ist typisch für Mahlos Sichtweise und läßt sich als Reflexionsprinzip formulieren: Das, was ganz am Ende der Ordinalzahlreihe stattfinden würde, könnten wir sie noch verlängern, findet bereits irgendwo auf der Ordinalzahlreihe statt.

Offenbar ist eine Mahlo-Kardinalzahl unerreichbar und ein Limes von unerreichbaren Kardinalzahlen, also 2-unerreichbar. „Stationär oft“ liegt insgesamt hinter allen durch Diagonalisierung von unten entstehenden Ausdünnungshierarchien der Unerreichbarkeit. Zwei Beispiele gibt der folgende Satz:

Satz (*Transzendenz von Mahlo-Kardinalzahlen über Unerreichbarkeits-Hierarchien*)

Sei κ eine Mahlo-Kardinalzahl. Dann ist κ eine κ -unerreichbare Kardinalzahl. Stärker gilt: $\{\mu < \kappa \mid \mu \text{ ist } \mu\text{-unerreichbar}\}$ ist stationär in κ .

Beweis

Sei $S = \{\alpha < \kappa \mid \alpha \text{ ist unerreichbar}\}$.

Wir definieren für $1 \leq \alpha < \kappa$ club-Mengen $C_\alpha \subseteq W(\kappa)$ rekursiv durch:

$$C_1 = S \cup \text{Lim}(S),$$

$$C_{\alpha+1} = \text{Lim}(C_\alpha \cap S) \quad \text{für } \alpha < \kappa,$$

$$C_\lambda = \bigcap_{\alpha < \lambda} C_\alpha \quad \text{für Limesordinalzahlen } \lambda < \kappa.$$

Für $1 \leq \alpha < \kappa$ ist dann $S_\alpha = C_\alpha \cap S$ stationär in κ .

Eine einfache Induktion über $1 \leq \alpha < \kappa$ zeigt:

$$S_\alpha = \{\mu < \kappa \mid \mu \text{ ist } \alpha\text{-unerreichbar}\} \text{ für alle } 1 \leq \alpha < \kappa.$$

Also ist κ eine κ -unerreichbare Kardinalzahl, denn $\kappa = \sup(S_\alpha)$ für alle $\alpha < \kappa$.

Sei $D = \Delta_{\alpha < \kappa} C_\alpha$. Dann ist D club in κ , also ist $S^* = D \cap S$ stationär in κ .

– Aber $S^* = \{\mu < \kappa \mid \mu \text{ ist } \mu\text{-unerreichbar}\}$.

Man kann nun analog zu α -unerreichbar α -Mahlo definieren, etwa:

„ κ ist 2-Mahlo“ falls „ $\{\alpha < \kappa \mid \alpha \text{ ist Mahlo}\}$ ist stationär in κ “,

„ κ ist 3-Mahlo“ falls „ $\{\alpha < \kappa \mid \alpha \text{ ist 2-Mahlo}\}$ ist stationär in κ “,

usw., bis hin zu

„ κ ist κ -Mahlo“ falls „ $\{\alpha < \kappa \mid \alpha \text{ ist } \beta\text{-Mahlo}\}$ ist stationär in κ für alle $\beta < \kappa$ “.

Die nächste Stufe wäre dann „ $\{\alpha < \kappa \mid \alpha \text{ ist } \alpha\text{-Mahlo}\}$ ist stationär in κ “, usw.

Über derartige Hierarchien hinausgehend interessiert sich der faustische Geist aber vor allem für Prinzipien, die andere Prinzipien transzendieren, so wie Mahlo-Kardinalzahlen die Hierarchien der Unerreichbarkeit transzendieren. Es gibt Dutzende solcher Prinzipien mit klingenden Namen wie *schwach kompakt*, *unbeschreibbar*, *subtil*, *meßbar*, *stark*, *superkompakt*, usw. Gerade noch in den Zeitrahmen 1870 – 1930, der den Kern dieses Buches bildet, fallen die meßbaren Kardinalzahlen. Der Leser mag ohnehin bereits fragen: Wo sind eigentlich die Ultrafilter geblieben? Eigenschaft (+++) von „großen Teilmengen“ sah nach einer vernünftigen Forderung aus ...

Meßbare Kardinalzahlen

Wir haben in den club-Filtern κ -vollständige Filter auf regulären $\kappa \geq \omega_1$ gefunden. $\mathcal{C}_\kappa$ ist aber niemals ein Ultrafilter: Es gibt $X \subseteq W(\kappa)$, die weder in $\mathcal{C}_\kappa$ noch in $\mathcal{F}_\kappa$ sind.

Triviale Beispiele für vollständige Ultrafilter sind nun:

Definition (*triviale Ultrafilter oder Hauptultrafilter*)

Sei M eine Menge und sei $x \in M$.

Dann heißt $U_x = \{X \subseteq M \mid x \in X\}$ der von x erzeugte triviale Ultrafilter oder Hauptultrafilter auf M .

Hauptultrafilter auf Mengen M mit $|M| = \kappa \geq \omega$ sind κ -vollständig. Andere Ultrafilter fallen einem nicht so ohne weiteres ein, und die Fragen formulieren sich nun fast von selbst:

Gibt es überhaupt nichttriviale Ultrafilter auf $W(\kappa)$?

Wenn ja: Gibt es κ -vollständige nichttriviale Ultrafilter auf $W(\kappa)$?

Die erste Frage läßt sich mit *ja* beantworten. Ulam hatte 1929 einen nichttrivialen Ultrafilter auf $W(\omega)$ konstruiert, und Tarski bewies dann das folgende allgemeine Resultat (1929):

Satz (*Existenzsatz von Tarski-Ulam über Ultrafilter*)

Sei M eine nichtleere Menge, und sei F ein Filter auf M .

Dann existiert ein Ultrafilter U auf M mit $U \supseteq F$.

Beweis

Sei $\mathcal{F} = \{G \supseteq F \mid G \text{ ist ein Filter auf } M\}$.

Dann ist $\mathcal{F}$ ein Zermelosystem (!). Sei also U ein Ziel von $\mathcal{F}$.

Wir zeigen, daß U ein Ultrafilter auf M ist.

Wegen $U \in \mathcal{F}$ ist U ein Filter.

Sei also $Y \subseteq M$. *Annahme* $Y \notin U$ und $M - Y \notin U$. Dann gilt:

(+) $X \cap Y \neq \emptyset$ für alle $X \in U$.

[Denn andernfalls existiert ein $X \in U$ mit $X \subseteq M - Y$, und dann wäre $M - Y \in U$ wegen $X \in U$ und U Filter.]

Wir setzen nun:

$U' = \{Z \mid Z \supseteq X \cap Y \text{ für ein } X \in U\}$.

Dann ist U' ein Filter auf M .

[$\emptyset \notin U'$ folgt aus (+).]

Ist $Z \in U'$ und $Z' \supseteq Z$, so ist $Z' \in U'$ nach Definition von U' .

Seien also $Z_1, Z_2 \in U'$. Dann existieren $X_1, X_2 \in U$ mit

$Z_1 \supseteq X_1 \cap Y$ und $Z_2 \supseteq X_2 \cap Y$.

Dann ist $Z_1 \cap Z_2 \supseteq X_1 \cap X_2 \cap Y$.

Also ist $Z_1 \cap Z_2 \in U'$ wegen $X_1 \cap X_2 \in U$, da U Filter.]

Offenbar gilt $U \subseteq U'$. Aber $Y = M \cap Y \in U'$, also $U \subset U'$,

– *im Widerspruch* zu U Ziel von $\mathcal{L}$.

Eine alternative Konstruktion ist: Wir zählen $\mathcal{P}(M) - F$ als $\langle X_\alpha \mid \alpha < \beta \rangle$ und fügen dann rekursiv zu F diejenigen X_α hinzu, für die $(+)$ gilt (wobei dann $U = F_\alpha$ jeweils der Stand der rekursiven Konstruktion ist). Das Endprodukt dieser Rekursion ist ein Ultrafilter auf M , der F fortsetzt.

Damit haben wir die Existenz nichttrivialer Ultrafilter auf unendlichen Mengen bewiesen: Ist $F = \{X \subseteq M \mid M - X \text{ endlich}\}$ und $U \supseteq F$ ein Ultrafilter, so ist U nichttrivial. Stärker wissen wir, daß es für alle regulären $\kappa \geq \omega_1$ einen Ultrafilter U gibt, der $\mathcal{C}_\kappa$ fortsetzt. Nichts in der Konstruktion dieses Ultrafilters garantiert uns aber, daß die κ -Vollständigkeit von $\mathcal{C}_\kappa$ bei der Erweiterung erhalten bleibt. In der Tat ist die Existenz eines nichttrivialen κ -vollständigen Ultrafilters ein starkes großes Kardinalzahlaxiom:

Definition (*meßbare Kardinalzahlen*)

Sei $\kappa \geq \omega_1$ eine Kardinalzahl. κ heißt *meßbar*, falls ein nichttrivialer κ -vollständiger Ultrafilter auf $W(\kappa)$ existiert.

Meßbare Kardinalzahlen gehen auf Ulam (1930) zurück. Ulams Untersuchungen konzentrierten sich auf Grundlagenprobleme der Maßtheorie, und seine meßbaren Kardinalzahlen heißen heute reellwertig meßbare Kardinalzahlen. Wir diskutieren Ulams Begriffe unten und verweilen erst einmal bei der obigen Definition der Meßbarkeit, die sich als die richtige herausgestellt hat.

Eine einfache Beobachtung ist:

Übung

Sei $\kappa \geq \omega_1$ regulär, und sei U ein Ultrafilter auf $W(\kappa)$.

Weiter sei $\tau \leq \kappa$ eine Kardinalzahl, $\tau \geq \omega_1$. Dann sind äquivalent.

(i) U ist τ -vollständig.

(ii) Für alle $\langle X_\alpha \mid \alpha < \lambda \rangle$, $\lambda < \tau$, gilt:

Ist $\bigcup_{\alpha < \lambda} X_\alpha \in U$, so existiert ein $\alpha < \lambda$ mit $X_\alpha \in U$.

[Das duale Primideal $I = \{X \subseteq W(\kappa) \mid W(\kappa) - X \in U\}$ hat die gleichen Vollständigkeitseigenschaften für Vereinigungen wie U für Schnitte.]

Weiter sind meßbare Kardinalzahlen in jedem Falle regulär:

Übung

Sei κ eine meßbare Kardinalzahl. Dann ist κ regulär.

[Andernfalls ist $W(\kappa) = \bigcup_{\alpha < \lambda} X_\alpha$ mit $\lambda < \kappa$ und $|X_\alpha| < \kappa$ für $\alpha < \lambda$.

Aber $X \notin U$ für $X \subseteq W(\kappa)$ mit $|X| < \kappa$ nach der Übung oben,

denn $X = \bigcup_{\alpha \in X} \{\alpha\}$ und $\{\alpha\} \notin U$ für alle $\alpha < \kappa$.]

Zwei Fragen, die sich stellen, sind:

- (1) Wie groß sind meßbare Kardinalzahlen?
- (2) Existieren ω_1 -vollständige Filter auf $W(\kappa)$ für gewisse reguläre $\kappa \geq \omega_1$?

Zur ersten Frage:

Satz (*Unerreichbarkeit von meßbaren Kardinalzahlen*)

Sei κ eine meßbare Kardinalzahl. Dann ist κ unerreichbar.

Beweis

Als meßbare Kardinalzahl ist κ regulär und überabzählbar.

Es genügt zu zeigen: $2^\mu < \kappa$ für alle $\mu < \kappa$ (dann ist κ ein starker Limes).

Annahme $2^\mu \geq \kappa$ für eine Kardinalzahl $\mu < \kappa$.

Dann existiert eine Folge $\langle Y_\alpha \mid \alpha < \kappa \rangle$ von paarweise verschiedenen Teilmengen von $W(\mu)$.

Sei U ein nichttrivialer κ -vollständiger Ultrafilter auf $W(\kappa)$.

Für $v < \mu$ setzen wir:

$$X_v = \begin{cases} \{\gamma < \kappa \mid v \in Y_\gamma\}, & \text{falls diese Menge in } U \text{ ist,} \\ \{\gamma < \kappa \mid v \notin Y_\gamma\}, & \text{sonst.} \end{cases}$$

Dann ist $X = \bigcap_{v < \mu} X_v \in U$.

Seien $\alpha, \beta \in X$. Dann gilt für alle $v < \mu$:

(+) $v \in Y_\alpha$ gdw $v \in Y_\beta$.

[Andernfalls ist o. E. $v \in Y_\alpha$ und $v \notin Y_\beta$.

Ist $\{\gamma < \kappa \mid v \in Y_\gamma\} \in U$, so ist $\beta \notin X_v$, also $\beta \notin X$, *Widerspruch*.

Ist $\{\gamma < \kappa \mid v \notin Y_\gamma\} \in U$, so ist $\alpha \notin X_v$, also $\alpha \notin X$, *Widerspruch*.]

— Also ist $\alpha = \beta$. Dann hat aber $X \in U$ höchstens ein Element, *Widerspruch*.

Die sich anschließende Frage, ob die kleinste meßbare Kardinalzahl echt größer als die kleinste unerreichbare Kardinalzahl ist, war mehrere Jahrzehnte lang ein offenes Problem, bis sie von Hanf und Tarski gegen 1960 über die Zwischenstufe der schwach kompakten Kardinalzahlen bejahend gelöst werden konnte [Hanf 1964]. Der Beweis ist heute modulo der in den 60er Jahren des 20. Jahrhunderts entwickelten Ultraprodukt-Technologie sehr einfach zu führen, sprengt aber den Rahmen dieser Darstellung. Es zeigt sich, daß sich meßbare Kardinalzahlen zu unerreichbaren Kardinalzahlen wie Elefanten zu Mücken verhalten. Meßbare κ sind insbesondere Mahlo-Kardinalzahlen. Man kann weiter zeigen, daß auf meßbaren Kardinalzahlen immer auch ein normaler nichttrivialer Ultrafilter existiert, und für einen normalen Ultrafilter U auf $W(\kappa)$ gilt dann z. B.:

$\{\alpha < \kappa \mid \alpha \text{ ist eine Mahlo-Kardinalzahl}\} \in U$,

d.h. aus der Sicht des Ultrafilters U besteht $W(\kappa)$ im wesentlichen aus Mahlo-Kardinalzahlen!

Der Leser kann versuchen, einige Hierarchien aus den genannten Kardinalzahlen zusammenzubauen und als grobe Zwischenstufen einer sehr feinen Hierarchie etwa betrachten:

„es gibt ein meßbares κ und eine unerreichbare Kardinalzahl größer als κ “,
 „es gibt ein meßbares κ und eine Mahlo-Kardinalzahl größer als κ “,
 „es gibt zwei meßbare Kardinalzahlen“,
 „es gibt einen unerreichbaren Limes κ von meßbaren Kardinalzahlen“,
 „es gibt ein κ mit $\{ \mu < \kappa \mid \mu \text{ ist meßbar } \}$ ist stationär in κ “,
 „es gibt einen meßbaren Limes κ von meßbaren Kardinalzahlen“, usw.

Meßbare Kardinalzahlen sind in der Mengenlehre unter anderem deswegen so prominent, weil für sie eine Umformulierung existiert, die in vielen Varianten dann immer wieder auftritt und den Kern der großen Kardinalzahlaxiome ans Licht bringt. Sie betrifft sog. elementare Einbettungen $j: V \rightarrow M$ (ungleich der Identität, M transitiv). Grob gesagt bedeutet das: Eine Aussage über ein Objekt x ist im Universum V genau dann wahr, falls $j(x)$ in M wahr ist. Ist j eine solche Einbettung, so existiert $\text{crit}(j) = \text{„das kleinste } \alpha \text{ mit } j(\alpha) > \alpha\text{“}$, der *kritische Punkt von j* . Die Existenz einer elementaren Einbettung mit $\kappa = \text{crit}(j)$ ist äquivalent zur Meßbarkeit von κ : Zur Konstruktion von j aus einem κ -vollständigen nichttrivialen Ultrafilter werden die erwähnten Ultraprodukte verwendet. Umgekehrt erhält man einen sogar normalen Ultrafilter U auf $W(\kappa)$ aus j mit $\text{crit}(j) = \kappa$ einfach durch $U = \{ X \subseteq W(\kappa) \mid \kappa \in j(X) \}$.

Noch stärkere Axiome fordern, daß sich M und V in $j: V \rightarrow M$ sehr ähnlich sind. Kenneth Kunen zeigte 1971, daß $M = V$ nicht gelten kann, und diese Unmöglichkeit einer elementaren Einbettung $j: V \rightarrow V$ ist vielleicht *die* obere (inkonsistente) Schranke für große Kardinalzahlaxiome.

Wir erwähnen noch ein weiteres fundamentales Resultat, das der Zusammenhang zwischen meßbaren Kardinalzahlen und elementaren Einbettungen in wenigen Zeilen liefert: In Gödels Modell L existieren keine meßbaren Kardinalzahlen. Dies hat Dana Scott 1961 gezeigt. Der Grund in einem Satz ist: Ist $j: L \rightarrow M$ eine elementare Einbettung, so gilt $M = L$.

Die zweite obige Frage betrifft das Verhältnis der Abschwächung der κ -Vollständigkeit zur ω_1 -Vollständigkeit. Hierzu:

Definition (*Additivität eines Filters*)

Sei M eine unendliche Menge, und sei F ein Filter auf M .

Dann ist die *Additivität von F* , in Zeichen $\text{Add}(F)$, definiert durch:

$\text{Add}(F) = \text{„das kleinste } \mu \text{ mit: Es gibt } X_\alpha \in F, \alpha < \mu, \text{ mit } \bigcap_{\alpha < \mu} X_\alpha \notin F\text{“}$,
 falls ein solches μ existiert.

Andernfalls setzen wir $\text{Add}(F) = |M|^+$.

Es gilt also $\text{Add}(F) \geq \omega$ für alle Filter F , und $\text{Add}(U) = \kappa$ für κ -vollständige nichttriviale Ultrafilter U auf $W(\kappa)$. Für triviale Ultrafilter U auf $W(\kappa)$ ist $\text{Add}(U) = \kappa^+$.

Die Additivität eines Filters ist immer eine reguläre Kardinalzahl:

Übung

Sei F ein Filter auf M . Dann ist $\text{Add}(F)$ regulär.

[Ist $X_\alpha \in F$ für $\alpha < \mu = \text{Add}(F)$, so ist $Y_\beta = \bigcap_{\alpha \leq \beta} X_\alpha \in F$ für $\beta < \mu$, und es gilt $\bigcap_{\alpha < \mu} X_\alpha = \bigcap_{\gamma < \text{cf}(\mu)} Y_{\beta_\gamma}$ für alle in μ konfinalen $\langle \beta_\gamma \mid \gamma < \text{cf}(\mu) \rangle$.]

Für Ultrafilter gilt nun stärker:

Satz (*Additivität und Meßbarkeit*)

Sei $\kappa \geq \omega_1$ eine Kardinalzahl, und sei U ein nichttrivialer Ultrafilter auf $W(\kappa)$ mit $\text{Add}(U) \geq \omega_1$.

Dann ist $\text{Add}(U)$ eine meßbare Kardinalzahl.

Beweis

Sei $\mu = \text{Add}(U)$. Dann existieren (paarweise disjunkte) $X_v \subseteq W(\kappa)$ mit $X_v \notin U$ für $v < \mu$ und $\bigcup_{v < \mu} X_v \in U$.

Wir definieren $E \subseteq \mathcal{P}(W(\mu))$ durch:

$Y \in E \text{ gdw } \bigcup_{v \in Y} X_v \in U$.

Offenbar ist E ein nichttrivialer Ultrafilter auf $W(\mu)$.

Weiter ist E μ -vollständig. Denn *andernfalls* gibt es ein $\lambda < \mu$ und $Y_\beta \subseteq W(\mu)$ mit $Y_\beta \notin E$ für alle $\beta < \lambda$ derart, daß $Y = \bigcup_{\beta < \lambda} Y_\beta \in E$.

Für $\beta < \lambda$ sei

$Z_\beta = \bigcup_{v \in Y_\beta} X_v$.

Dann ist $\bigcup_{\beta < \lambda} Z_\beta = \bigcup_{v \in Y} X_v \in U$.

Also existiert wegen $\lambda < \mu = \text{Add}(U)$ ein $\beta^* < \lambda$ mit $Z_{\beta^*} \in U$.

– Dann ist aber $Y_{\beta^*} \in E$ nach Definition von E , *Widerspruch*.

Korollar

Sei κ die kleinste Kardinalzahl, für die ein ω_1 -vollständiger nichttrivialer Ultrafilter auf $W(\kappa)$ existiert. Dann ist κ meßbar.

Schon die Existenz eines nichttrivialen Ultrafilters auf $W(\kappa)$, der gegenüber abzählbaren Durchschnitten abgeschlossen ist, impliziert also die Existenz einer meßbaren Kardinalzahl $\mu \leq \kappa$!

Maße auf überabzählbaren Mengen

Das Folgende würde eine ausführlichere separate Darstellung mit einer sorgfältigen und historisch orientierten Einleitung verdienen. Wir begnügen uns hier mit einem Überblick, der viele ad-hoc-Definitionen enthalten wird. Ziel ist lediglich, die Brücke zur Maßtheorie aufzuzeigen, und dies nicht nur aus historischer Sicht. Der nicht an Fragen der reellwertigen Maßtheorie interessierte Leser kann diesen Zwischenabschnitt problemlos überschlagen.

Die Frage nach der Existenz von vollständigen Ultrafiltern ergab sich zwanglos aus dem Studium des Filterbegriffs. Historisch verlief die Sache auf einem Nebengleis, dem die Theorie den Namen „meßbar“ zu verdanken hat: Ulam untersuchte Fragen der Maßtheorie, als er auf das Phänomen hinter den meßbaren Kardinalzahlen stieß.

Mathematiker lieben die abstrakte Meßkunst. Sie ordnen dabei einer Menge eine nichtnegative reelle Zahl zu – das Maß der Menge. „Maß von“ ist wie „große Teilmenge von“ extensional „in einem bestimmten Sinne“ zu fassen, in diesem Fall durch eine Funktion. Die zugehörigen mathematischen Begriffe kristallisierten sich Anfang des 20. Jahrhunderts heraus. Wir verwenden sie hier in der folgenden Form:

Definition (*Inhalt*)

Sei M eine unendliche Menge, und sei $I : \mathcal{P}(M) \rightarrow [0, 1] \subseteq \mathbb{R}$.

I heißt ein (*nichttrivialer*) *Inhalt auf M* , falls gilt:

- (i) $I(\emptyset) = 0$, $I(M) = 1$, (*Normierung*)
- (ii) $I(\{x\}) = 0$ für alle $x \in M$, (*Nichttrivialität*)
- (iii) $I(A) \leq I(B)$ für alle $A \subseteq B \subseteq M$, (*Monotonie*)
- (iv) $I(\bigcup_{0 \leq i \leq n} A_i) = \sum_{0 \leq i \leq n} I(A_i)$ für alle $n < \omega$ und alle paarweise disjunkten $A_i \subseteq M$, $0 \leq i \leq n$. (*endliche Additivität*)

Definition (*Maß*)

Sei I ein Inhalt auf M . I heißt ein (*nichttriviales*) *Maß auf M* , falls zusätzlich folgende σ -Additivität (oder ω_1 -Additivität) gilt:

$$I\left(\bigcup_{i < \omega} A_i\right) = \sum_{i < \omega} I(A_i) \quad (= \sup_{n < \omega} \sum_{0 \leq i \leq n} I(A_i))$$

für alle paarweise disjunkten $A_i \subseteq M$, $i < \omega$.

Der bevorzugte Buchstabe für Maße ist im folgenden nicht I , sondern μ .

Analytiker, Wahrscheinlichkeitstheoretiker und Physiker wären ohne abzählbare Grenzübergänge und Vertauschungen von bestimmten Operationen – etwa das Herausziehen eines Limes aus einem Integral – höchst unglückliche Geschöpfe, und dies erklärt, warum in vielen Bereichen der Mathematik von Maßen die Rede ist, und Inhalte eine eher untergeordnete Rolle spielen. Die σ -Additivität wurde bereits von Borel 1898 in den Mittelpunkt gerückt.

Die nichttrivialen Ultrafilter auf M entsprechen den 0-1-wertigen Inhalten auf M , d. h. den Inhalten I mit $\text{rng}(I) = \{0, 1\}$: Gegeben U betrachten wir den Inhalt I mit $I(A) = 1$, falls $A \in U$, $I(A) = 0$ sonst. Gegeben einen 0-1-wertigen Inhalt I setzen wir $U = \{A \subseteq M \mid I(A) = 1\}$. Nach dem Ultrafiltersatz von Ulam-Tarski existiert insbesondere also immer ein Inhalt auf einer unendlichen Menge M .

Analog entsprechen die ω_1 -vollständigen nichttrivialen Ultrafilter auf M den 0-1-wertigen Maßen auf M . Wir wissen, daß vollständige Ultrafilter schwer zu haben sind. Im Gegensatz zur Ultrafiltertheorie haben wir nun aber die übliche Flut reeller Werte zur Verfügung, und nicht nur 0 oder 1. Maße werden also vielleicht leichter zu konstruieren sein ...

Maße auf abzählbaren Mengen lassen sich leicht angeben: Für $A \subseteq W(\omega)$ setze $\mu(A) = \sum_{n \in A} 1/2^{(n+1)}$. Dann ist μ ein Maß auf $W(\omega)$ ($1/2 + 1/4 + 1/8 + \dots = 1$). Dagegen existieren keine 0-1-wertigen Maße auf $W(\omega)$!

Zwei Grundprobleme der Maßtheorie lauten:

- (1) *Existieren Maße auf überabzählbaren Mengen?*
- (2) *Gibt es ein Maß auf $\mathbb{R}$ oder auf einem beschränkten Intervall von $\mathbb{R}$?*

Die erste Idee für ein Maß auf z. B. $I = [0, 1]$ ist: $\mu(A)$ = „das Längenmaß von A “ für $A \subseteq I$. Dies sollte doch für alle $A \subseteq I$ in vernünftiger Weise zu definieren sein. Es stellt sich heraus, daß hier große Schwierigkeiten auftreten, und daß diese Schwierigkeiten ganz allgemein sind und nichts mit einem speziellen Längenmaß zu tun haben.

Die Frage, ob für alle $A \subseteq I$ ein Längenmaß erklärt werden kann, hatte Henri Lebesgue 1902 gestellt. Der Begriff „Längenmaß“ wird durch eine zusätzliche Bedingung eingefangen:

(+) Für alle $A \subseteq I$ und $0 \leq x \leq 1$ ist $\mu(A + x) = \mu(A)$. (*Translationsinvarianz*)

Hier ist (der Bequemlichkeit halber) $I = [0, 1[$ und $A + x = \{y + x \mid y \in A, y + x < 1\} \cup \{y + x - 1 \mid y \in A, y + x \geq 1\}$ die Rechts-Translation von A um x modulo 1, d. h. Punkte, die das Intervall bei der Verschiebung um x rechts verlassen, tauchen links wieder auf. Man kann sich I auch als Kreislinie vorstellen, indem 0 und 1 verklebt werden. $A + x$ ist dann die Drehung von A um den Winkel $x \cdot 2\pi$.

Vitali hatte 1905 gezeigt, daß die Translationsinvarianz eine zu starke Forderung an Maße auf I darstellt. Unabhängig wurde Vitalis Gegenbeispiel auch von Hausdorff gefunden. Wir setzen für $x, y \in I$ (vgl. 1.3):

$x \sim y$ falls $|x - y| \in \mathbb{Q}$.

Dann ist $\sim$ eine Äquivalenzrelation auf I . Es gilt:

Übung (Vitali-Hausdorff)

Es gibt kein Maß auf $I = [0, 1[$ mit (+).

[Sei $A \subseteq I$ ein vollständiges Repräsentantensystem für $\sim$. Dann gilt $I = \bigcup_{q \in \mathbb{Q}} A + q$, und die Mengen $A + q$, $q \in \mathbb{Q}$ sind paarweise disjunkt.

Ist $\mu(A) = 0$, so wäre $1 = \mu(I) = 0 + 0 + \dots = 0$, *Widerspruch*.

Ist $\mu(A) = \varepsilon > 0$, so wäre $1 = \varepsilon + \varepsilon + \dots$, *Widerspruch*, denn für alle $\varepsilon > 0$ existiert ein $n < \omega$ mit $\varepsilon \cdot n > 1$.]

Analoges gilt für $\mathbb{R}$ statt I , wobei in einer Maßtheorie für ganz $\mathbb{R}$ ein Maß auf $\mathbb{R}$ dann nur noch σ -finit sein muß, d. h.: Die Maße von beschränkten Teilmengen von $\mathbb{R}$ sind endlich. (Ein Maß auf $\mathbb{R}$ ist dann eine Funktion $\mu: \mathcal{P}(\mathbb{R}) \rightarrow [0, \infty[\cup \{\infty\}$.) Auch hier führt die Vitali-Hausdorff-Konstruktion zu einem Widerspruch. Die Translation ist dann sogar die echte Translation, $A + x = \{y + x \mid y \in A\}$. Wie in der Übung ist $\mu(A) = 0$ für ein vollständiges Repräsentantensystem von $\sim$ nicht möglich. Sei also $\mu(A) = \varepsilon > 0$, und sei $\varepsilon \cdot n > \mu([0, 2]) = z$. O. E. dürfen wir $A \subseteq [0, 1]$ annehmen. Für paarweise verschiedene rationale $0 < q_1, \dots, q_n < 1$ ist dann $z < \sum_{1 \leq i \leq n} \mu(A + q_i) \leq z$ wegen $\bigcup_{1 \leq i \leq n} A + q_i \subseteq [0, 2]$. Widerspruch!

Volle Translationsinvarianz ist also nicht möglich. Sie erzeugt zuviel Symmetrie, und diese verletzt die abzählbare Additivität.

Symmetrie ist auch für Inhalte eine starke Forderung. Es gibt translationsinvariante bzw. (für höhere Dimensionen) kongruenzinvariante Inhalte auf $\mathbb{R}^1$ und $\mathbb{R}^2$ [Banach 1923], nicht aber auf $\mathbb{R}^n$ für $n \geq 3$ [Hausdorff 1914].

Stefan Banach (1892 – 1945) stellte dann die Frage nach der Existenz eines *beliebigen* nichttrivialen Maßes auf $I = [0, 1]$. Diese Frage ist wesentlich subtiler.

1929 zeigten Banach und Kuratowski, daß unter Annahme der Kontinuumshypothese kein Maß auf $I = [0, 1]$ existiert. Im Jahr 1930 gab Ulam einen weiteren Beweis. In der Tat folgt der Satz leicht aus der Existenz von Ulam-Matrizen auf ω_1 und der folgenden einfachen Beobachtung:

Übung

Sei μ ein Maß auf M , und sei $\mathcal{A} \subseteq \mathcal{P}(M)$ eine Menge von paarweise disjunkten Teilmengen von M mit $\mu(A) > 0$ für alle $A \in \mathcal{A}$. Dann ist $\mathcal{A}$ abzählbar.

[Sei $\mathcal{A}_n = \{A \in \mathcal{A} \mid \mu(A) > 1/n\}$ für $1 \leq n < \omega$. Dann ist jedes $\mathcal{A}_n$ sogar endlich. Aber $\mathcal{A} = \bigcup_{1 \leq n < \omega} \mathcal{A}_n$, also ist $\mathcal{A}$ abzählbar.]

Damit haben wir:

Satz (Existenz von Maßen und die Kontinuumshypothese)

Es gelte (CH). Dann existiert kein Maß auf $I = [0, 1] \subseteq \mathbb{R}$.

Beweis

Annahme doch. Nach Voraussetzung gilt $|I| = |W(\omega_1)|$.

Für beliebige Mengen M, N übersetzt ein bijektives $f: M \rightarrow N$ ein Maß μ auf M in ein Maß ν auf N via

$$\nu(A) = \mu(f^{-1}A) \quad \text{für } A \subseteq N.$$

Nach Annahme existiert also ein Maß μ auf $W(\omega_1)$.

Sei $\langle X_{n,\alpha} \mid n < \omega, \alpha < \omega_1 \rangle$ eine Ulam-Matrix auf ω_1 .

Für alle $\alpha < \omega_1$ existiert dann ein $n < \omega$ mit $\mu(X_{n,\alpha}) > 0$.

$$[\text{Denn } \mu(\bigcup_{n < \omega} X_{n,\alpha}) = \mu(W(\omega_1) - W(\alpha)) = 1 \text{ für alle } \alpha < \omega_1.]$$

Wie früher existieren dann ein $n^* < \omega$ und ein $A \subseteq \omega_1$ mit $|A| = \omega_1$ und $\mu(X_{n^*,\alpha}) > 0$ für alle $\alpha \in A$. Also ist μ strikt positiv auf überabzählbar vielen

– paarweise disjunkten Mengen, *Widerspruch*.

Was ist, wenn $2^\omega = \omega_2$ gilt? Ist ein Maß auf $I = [0, 1]$ überhaupt möglich, oder ist die Existenz eines reellwertigen Maßes auf einer überabzählbaren Menge M doch gleichwertig mit der Existenz eines 0-1-wertigen Maßes auf M ? Dann wäre ein Maß auf einer Menge M mit $|M| = |\mathbb{R}|$ nicht möglich, denn Maße auf überabzählbaren Mengen würden dann nur auf meßbaren Kardinalzahlen existieren, und diese sind starke Limeskardinalzahlen, also viel größer als $|\mathbb{R}| = 2^\omega$.

Die Antwort ist:

Die Existenz eines Maßes auf $I = [0, 1]$ ist modulo der Existenz von meßbaren Kardinalzahlen widerspruchsfrei (Solovay 1966, 1971). Existiert ein Maß auf I , so existiert eine schwach unerreichbare Kardinalzahl κ mit $|\mathbb{R}| \geq \kappa$ (Ulam 1930).

Das Kontinuum muß also sehr groß sein, wenn ein Maß auf dem Einheitsintervall existiert. Es ist viel, viel größer als ω_1 , aber immer noch viel, viel kleiner als die erste stark unerreichbare Kardinalzahl oder gar die erste meßbare Kardinalzahl. Insgesamt ist die Frage nach der Existenz von Maßen auf dem Einheitsin-

tervall unabhängig, und darüber hinaus ist die Annahme der Existenz eines solchen Maßes eine starke Hypothese (de facto von der Stärke von „es gibt eine meßbare Kardinalzahl“, wie Solovay ebenfalls gezeigt hat). (Für „Einheitsintervall“ kann man hier jede Menge der Mächtigkeit von $\mathbb{R}$ einsetzen.)

Wir geben hier nur noch den Beweis, daß die Existenz eines Maßes auf I die Mächtigkeit des Kontinuums beträchtlich nach oben zieht.

Definition (*τ -additiv, Additivität eines Maßes*)

Sei μ ein Maß auf einer Menge M .

μ heißt *τ -additiv* oder *τ -vollständig* für eine Kardinalzahl τ , falls für alle $\lambda < \tau$ und alle paarweise disjunkten $X_\alpha \subseteq M$, $\alpha < \lambda$, gilt:

$$\mu\left(\bigcup_{\alpha < \lambda} X_\alpha\right) = \sum_{\alpha < \lambda} \mu(X_\alpha) \quad (= \sup_{E \subseteq W(\lambda), E \text{ endlich}} \sum_{\alpha \in E} \mu(X_\alpha)).$$

Wir setzen $\text{Add}(\mu) =$ „das kleinste τ mit: μ ist nicht τ^+ -additiv“.

$\text{Add}(\mu)$ heißt die *Additivität* des Maßes μ .

Die σ -Additivität eines Maßes μ ist identisch zur ω_1 -Additivität von μ nach dieser Definition. Es gilt also $\text{Add}(\mu) \geq \omega_1$ für alle Maße μ . Ist $\text{Add}(\mu) = \tau$, so existieren paarweise disjunkte $X_\alpha \subseteq M$, $\alpha < \tau$, mit $\mu(\bigcup_{\alpha < \tau} X_\alpha) > \sum_{\alpha < \tau} \mu(X_\alpha)$, und τ ist minimal hierfür. Für alle Maße μ auf M mit $|M| = \kappa$ gilt $\text{Add}(\mu) \leq \kappa$, denn es gilt $1 = \mu(M) > \sum_{x \in M} \mu(\{x\}) = \sup_{E \subseteq M, E \text{ endlich}} \sum_{x \in E} \mu(\{x\}) = 0$.

Definition (*reellwertig meßbare Kardinalzahl*)

Sei $\kappa \geq \omega_1$ eine Kardinalzahl. κ heißt *reellwertig meßbar*, falls ein κ -additives Maß auf $W(\kappa)$ existiert.

Satz (*reellwertig meßbare Kardinalzahlen sind schwach unerreichbar*)

Sei κ reellwertig meßbar. Dann ist κ schwach unerreichbar.

Beweis

κ ist regulär:

Wie für meßbare Kardinalzahlen.

κ ist eine Limeskardinalzahl:

Annahme $\kappa = \mu^+$. Sei dann $\langle X_{v,\alpha} \mid v < \mu, \alpha < \kappa \rangle$ eine Ulam-Matrix auf κ .

Wegen der κ -Additivität von μ gibt es in jeder Spalte einen Eintrag mit positivem Maß. Dann existiert wieder eine Zeile mit überabzählbar vielen paarweise disjunkten Einträgen mit positivem Maß, *Widerspruch!*

Reellwertig meßbare Kardinalzahlen sind also schwach unerreichbar. Im Gegensatz zu meßbaren Kardinalzahlen sind sie jedoch i.a. nicht unerreichbar.

Analog zu den meßbaren Kardinalzahlen gilt:

Satz (*Additivität und reellwertige Meßbarkeit*)

Sei $\kappa \geq \omega_1$, und sei μ ein Maß auf $W(\kappa)$.

Dann ist $\text{Add}(\mu)$ reellwertig meßbar.

Beweis

Wir folgen dem Beweis des Satzes, daß $\text{Add}(U)$ eine meßbare Kardinalzahl ist für nichttriviale Ultrafilter U auf $W(\kappa)$ mit $\text{Add}(U) \geq \omega_1$.

Sei $\tau = \text{Add}(\mu) \geq \omega_1$. Dann existieren paarweise disjunkte $X_v \subseteq W(\kappa)$ mit $\mu(X_v) = 0$ für $v < \tau$ und $a = \mu(\bigcup_{v < \tau} X_v) > 0$.

[Es existieren τ -viele paarweise disjunkte Mengen $Y_v \subseteq W(\kappa)$ mit $b = \mu(\bigcup_{v < \tau} Y_v) > \sum_{v < \tau} \mu(Y_v)$. Sei $A = \{v \mid \mu(Y_v) > 0\}$.

Dann ist A abzählbar, und jede Aufzählung $\langle X_v \mid v < \tau \rangle$ von $\{Y_v \mid v \in W(\tau) - A\}$ ist wie gewünscht, denn $\mu(\bigcup_{v \in A} Y_v) < b$.]

Wir definieren $\mu' : \mathcal{P}(W(\tau)) \rightarrow [0, 1]$ durch:

$$\mu'(A) = \mu(\bigcup_{v \in A} X_v) / a \quad \text{für } A \subseteq W(\tau).$$

Dann ist μ' ein τ -additives Maß auf $W(\tau)$ (!).

— Also ist τ eine reellwertig meßbare Kardinalzahl.

Korollar

Sei κ die kleinste Kardinalzahl, für die ein Maß auf $W(\kappa)$ existiert.

Dann ist κ reellwertig meßbar.

Weiter gilt: Existiert ein Maß auf einer Menge M mit $|M| = |\mathbb{R}|$, so existiert ein schwach unerreichbares κ und es gilt $|\mathbb{R}| \geq \kappa$.

Die Maßtheorie reagiert auf dieses bedrohliche und irritierende Grundproblem der Existenz von Maßen mit einem taktischen Rückzug: Der Definitionsbereich einer Maßfunktion μ wird eingeschränkt. Maße sind weiterhin σ -additiv, sie sind aber nur noch auf einem Teilsystem von $\mathcal{P}(M)$ definiert, einer sog. σ -Algebra, und i. a. nicht mehr auf ganz $\mathcal{P}(M)$. Dies genügt für viele Anwendungen. ($\mathcal{A} \subseteq \mathcal{P}(M)$ heißt eine σ -Algebra, falls gilt: (i) $\emptyset \in \mathcal{A}$, (ii) $X \in \mathcal{A}$ folgt $M - X \in \mathcal{A}$ für alle X , (iii) $\bigcup_{n < \omega} A_n \in \mathcal{A}$ für alle $A_n \in \mathcal{A}$, $n < \omega$.) Für die reellen Zahlen sind dann (σ -finite) Maße zumeist auf der Borel- σ -Algebra $\mathfrak{B} = \mathfrak{B}(\mathbb{R})$ definiert (vgl. 2. 7). Speziell existiert ein eindeutiges translationsinvariantes Maß μ auf $\mathfrak{B}$ mit $\mu([a, b]) = b - a$ für reelle $a < b$, das sog. Lebesgue-Maß. Die übliche Definition des Lebesgue-Maßes liefert eine σ -Algebra, die größer ist als $\mathfrak{B}$. Die genaue Bestimmung dieser σ -Algebra der Lebesgue-meßbaren Mengen ist eine komplexe Angelegenheit, und viele Zugehörigkeitsfragen sind von der üblichen Mengenlehre unabhängig. Genügend große Kardinalzahlen implizieren die Lebesgue-Meßbarkeit aller projektiven Mengen (vgl. 2. 12).

Vorläufiges Schlußwort zu großen Kardinalzahlen

Bereits die möglichen Typen von wohlgeordneten linearen Ordnungen erweisen sich als unerwartet komplex. Man kann Fragen stellen wie „Gibt es einen unerreichbaren Typus?“, „Gibt es sogar einen Mahlo-Typus?“, „Gibt es sogar einen meßbaren Typus?“ die unbeantwortbar sind und in ihrem Wahrheitsgehalt erst einmal offen bleiben müssen. Das Schöne ist: Dem Neugierigen eröffnet die Mathematik ein ganz neuartiges Land voller Naturschönheiten, reich an Bodenschätzen, frei in der Verfassung. Niemand konnte den Reichtum, der auf einer platonisch betrachteten Ordinalzahlreihe herrschen mag, voraussehen. Ganz unabhängig von „wahr“ oder „falsch“ besteht Einigkeit darin, daß man eine

Struktur gefunden hat, die sich objektiv erforschen läßt. Entdeckt wurde eine wunderbare Welt außerhalb von uns, deren Existenz niemand ernsthaft bezweifeln kann. Gegenstand der Diskussion bleibt dann nur noch die Feinanalyse des Unterschiedes zwischen „ $\mathbb{R}$ existiert“ und „es gibt eine unerreichbare Kardinalzahl“, sowie zwischen „es gibt eine unerreichbare Kardinalzahl“ und „es gibt eine meßbare Kardinalzahl“. Gerade in dieser Feinanalyse scheiden sich die Geister. Eine Scheu vor großen Kardinalzahlaxiomen, wie sie bei einer ersten Begegnung zuweilen auftritt, hat aber wohl niemand mehr, der ein paar Abende mit ihnen verbracht hat. Man gewöhnt sich schnell daran, von irgendwo weit oben in den Ordinalzahlen auf eine Unzahl von unerreichbaren Kardinalzahlen hinabzublicken, anstatt bereits vor der ersten unerreichbaren Zahl von ω heraufblickend Angst zu bekommen. Und das gleiche gilt für „Mahlo“ und „meßbar“.

Wir haben hier gerade einmal das Portal in die Welt der starken Axiome durchschritten, und die eigentlichen Oasen, die dort zu entdecken sind, haben wir mit unserer elementar ausgestatteten Karawane gar nicht erreicht. Die Gesamtheit der Theorie erst ist es, die besticht. Dieser kurze Ausflug konnte und wollte niemanden davon überzeugen, daß derlei Dinge so wahnsinnig interessant und wichtig sein sollen. Die Mengenlehre selbst hat weit mehr als ein halbes Jahrhundert nach Mahlos und Ulams ersten Vorstößen gebraucht, um sich dieses Reich zu erschließen. Belassen wir es also hier bei diesem kleinen Ausblick. Er zeigt nicht, wie es dort draußen wirklich zugeht, und kann die Einheit in der Vielfalt, die dort herrscht, nicht wiedergeben. Aber neugierig machen wollte er schon.

Das Kapitel schließt mit Paul Mahlos unerschrockenem Existenzkriterium für große Kardinalzahlen: Ihre durch eine sorgfältige Untersuchung gestützte anscheinende Widerspruchsfreiheit und ihre fügsame Eingliederung in die Hierarchie der bereits untersuchten Zahlen bilden, so Mahlo, genügend Dokumente, um ihnen das Grundrecht der mathematischen Existenz verbrieften zu können.

Paul Mahlo über die Existenz großer Kardinalzahlen

„Da wir im folgenden die Existenz gewisser bisher unbekannter Transfiniten behaupten, zwingt uns der Widerspruch im Begriffe von W [der Klasse aller Ordinalzahlen, vgl. Kap. 13], die Bedingungen für die Existenz einer neu definierten Transfiniten möglichst vollständig anzugeben. Dies ist vor allem deshalb nötig, weil, wie bei den Anfangszahlen ω und Ω der zweiten und dritten Zahlenklasse [wie bei ω und ω_1] oder etwa der ersten regulären Anfangszahl mit einem Limeszahindex [der ersten schwach unerreichbaren Kardinalzahl]..., kein Widerspruch auffindbar war, aber jeder derartigen Neuschöpfung ein starkes Mißtrauen entgegengebracht wird. Wir schreiben deshalb jeder Transfiniten, in deren Definition wir keinen Widerspruch finden können, die Existenz zu, wenn sich ihre Gleichheit oder Verschiedenheit von jeder sonst definierten Transfiniten feststellen läßt. Die notwendige Bedingung ist auch hinreichend, weil sie die Rangordnung zwischen beliebigen Transfiniten liefert... Die obige Bedingung kann aber für die Praxis lange Zeit wertlos sein, indem sich etwa bei der Unbekanntheit gewisser Transfiniten ein Widerspruch in der Definition vermeintlicher Transfiniten verbirgt.“

(Paul Mahlo 1911, „Über lineare transfinite Mengen“)

10. Die Ordnungstypen von $\mathbb{Q}$ und $\mathbb{R}$

Wir untersuchen in diesem Kapitel den Ordnungstyp η der rationalen Zahlen, und den Ordnungstyp θ der reellen Zahlen. Unser Ziel ist, diese beiden Typen ordnungstheoretisch zu charakterisieren: wir wollen einfache Eigenschaften finden, welche genau von den zu den rationalen bzw. den reellen Zahlen ähnlichen linearen Ordnungen erfüllt werden.

Für den Ordnungstyp η der rationalen Zahlen gibt es eine erstaunlich einfache Charakterisierung, die uns überraschende neue Beweise der Existenz transzendenter Zahlen und der Überabzählbarkeit der reellen Zahlen liefern wird. Für diese Charakterisierung brauchen wir die bereits für die reelle Ordnung definierten Begriffe „unbeschränkt“ und „dicht“ allgemein für lineare Ordnungen.

Definition (*unbeschränkte lineare Ordnungen*)

Sei $\langle M, < \rangle$ eine lineare Ordnung.

- (i) $\langle M, < \rangle$ heißt *nach links unbeschränkt*, falls die Ordnung kein kleinstes Element hat, d. h. falls es für alle $x \in M$ ein $y \in M$ gibt mit $y < x$.
- (ii) $\langle M, < \rangle$ heißt *nach rechts unbeschränkt*, falls die Ordnung kein größtes Element hat, d. h. falls es für alle $x \in M$ ein $y \in M$ gibt mit $x < y$.
- (iii) $\langle M, < \rangle$ heißt *unbeschränkt*, falls $\langle M, < \rangle$ nach links und rechts unbeschränkt ist.

Definition (*dichte lineare Ordnungen*)

Eine lineare Ordnung $\langle M, < \rangle$ heißt *dicht*, falls für alle $x, y \in M$ mit $x < y$ ein $z \in M$ existiert mit $x < z < y$.

Offenbar ist jede dichte Ordnung unendlich.

Der Ordnungstyp η

Für die Ordnung der rationalen Zahlen $\langle \mathbb{Q}, < \rangle$ gilt:

- (i) $\langle \mathbb{Q}, < \rangle$ ist abzählbar,
- (ii) $\langle \mathbb{Q}, < \rangle$ ist unbeschränkt,
- (iii) $\langle \mathbb{Q}, < \rangle$ ist dicht.

Cantor hat nun gesehen, daß diese drei Bedingungen den Ordnungstyp η bereits charakterisieren:

Satz (*Cantors Isomorphiesatz für $\mathbb{Q}$*)

Sei $\langle M, < \rangle$ eine lineare Ordnung. Es gelte:

- (i) $\langle M, < \rangle$ ist abzählbar,
- (ii) $\langle M, < \rangle$ ist unbeschränkt,
- (iii) $\langle M, < \rangle$ ist dicht.

Dann gilt $\eta = \text{o. t.}(\langle M, < \rangle)$, d. h. $\langle M, < \rangle$ und $\langle \mathbb{Q}, < \rangle$ sind ähnlich.

Beweis

Seien $x_0, x_1, \dots$ und $q_0, q_1, \dots$ Aufzählungen von M bzw. $\mathbb{Q}$ (ohne Wiederholungen; beide Mengen sind abzählbar unendlich).

Wir nennen ein $g : M' \rightarrow \mathbb{Q}$, $M' \subseteq M$, ordnungserhaltend, falls für alle $x, y \in M'$ gilt:

$$x < y \text{ gdw } g(x) < g(y).$$

Wir definieren nun ordnungserhaltende $f_n : M_n \rightarrow \mathbb{Q}$, $M_n \subseteq M$, durch Rekursion über $n \in \mathbb{N}$, wobei wir in jedem Schritt f_n um zwei neue Elemente erweitern.

Sei also $n \in \mathbb{N}$ und f_{n-1} definiert, wobei $f_{-1} = \emptyset$.

Wir setzen

$$f_n = f_{n-1} \cup \{ (x_{k(n)}, q_{\ell(n)}), (x_{k'(n)}, q_{\ell'(n)}) \},$$

wobei:

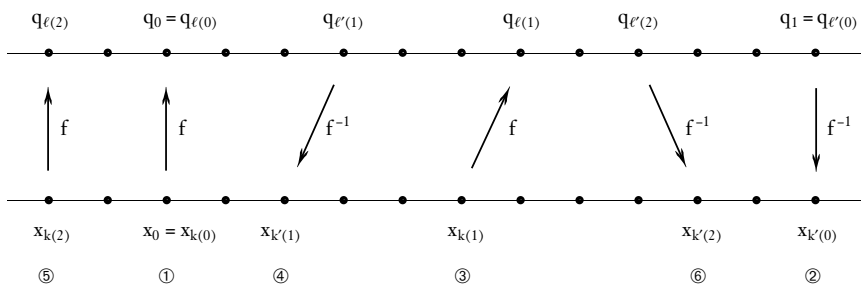
$k(n)$ = „das kleinste $k \in \mathbb{N}$ mit: $x_k \notin \text{dom}(f_{n-1})$ “,

$\ell(n)$ = „das kleinste $\ell \in \mathbb{N}$ mit: $f_{n-1} \cup \{ (x_{k(n)}, q_\ell) \}$ ist ordnungserhaltend“,

$\ell'(n)$ = „das kleinste $\ell \in \mathbb{N}$ mit: $q_\ell \notin \text{rng}(f_{n-1}) \cup \{ q_{\ell(n)} \}$ “,

$k'(n)$ = „das kleinste $k \in \mathbb{N}$ mit: $f_{n-1} \cup \{ (x_{k(n)}, q_{\ell(n)}), (x_k, q_{\ell'(n)}) \}$ ist ordnungserhaltend“.

$\ell(n)$ und $k'(n)$ existieren wegen der Dichtheit und Unbeschränktheit der Ordnungen.



Sei $f = \bigcup_{n \in \mathbb{N}} f_n$.

Nach Konstruktion ist $f : M \rightarrow \mathbb{Q}$ bijektiv und ordnungserhaltend.

– Somit $\eta = \text{o. t.}(\langle M, < \rangle)$.

Wir definieren hier f im „Zickzack“, wodurch der Nachweis der Surjektivität von f trivial wird. Der Trick ist als „back and forth“-Argument bekannt und auch an anderen Stellen nützlich (vor allem in der Modelltheorie). Das Argument geht auf Maurice Fréchet zurück (Fréchet, 1910).

Schoenflies (1913): „Die wichtigsten Ordnungstypen unendlicher Mengen sind diejenigen der Menge R aller rationalen, resp. der Menge L aller reellen Zahlen, falls sie der Größe nach geordnet werden. Den Ordnungstyp der Menge R aller rationalen Zahlen hat Cantor ... durch η bezeichnet und von ihm den folgenden wichtigen Satz bewiesen:

Alle abzählbaren überall dichten linearen Mengen ohne erstes und letztes Element haben den gleichen Ordnungstypus wie die rationalen Zahlen.

Ist nämlich M eine solche Menge, so wird der Beweis geführt sein, wenn wir eine Abbildung von M auf R so ausführen können, daß die Anordnung gewahrt bleibt. Wir setzen dazu die Menge R in die Form

$$R = \{ r_v \} = r_1, r_2, \dots, r_v, \dots,$$

..., und setzen analog M in die Form

$$M = \{ m'_v \} = m'_1, m'_2, \dots, m'_v, \dots$$

Wir ordnen zunächst dem r_1 das Element $m'_1 = m_1$ zu, und wählen als m_2 das Element m'_λ mit kleinstem Index, das zu m_1 ebenso liegt, wie r_2 zu r_1 . Als m_3 wählen wir analog das Element m'_μ mit kleinstem Index, das zu m_1 und m_2 ebenso liegt, wie r_3 zu r_1 und r_2 . So fahren wir fort, dann ist klar, daß wir zu jedem r_k ein und nur ein Element m_k erhalten; es ist nur noch zu zeigen, daß auch das Umgekehrte der Fall ist. Dies ergibt sich aufgrund einfacher Induktion ...

Man kann die vorstehende Abbildung auch in folgender Weise ausführen.²⁾ [Fußnote 2): Vgl. M. Fréchet, Math. Ann. 68 (1910), S. 151...] Man ordne wiederum die Punkte r_1 und m'_1 einander zu, und bezeichne sie nunmehr durch x_1 und y_1 . Jetzt suche man in R und M die ersten Punkte, die links und rechts von x_1 und y_1 liegen; sie seien in ursprünglicher Bezeichnung r_v und r_μ , resp. m'_λ und m'_κ ; in neuer Bezeichnung seien sie x_2, x_3 und y_2, y_3 . Analog suche man in R und M die ersten Punkte, die in die durch x_1, x_2, x_3 und y_1, y_2, y_3 bestimmten Intervalle fallen; sie mögen durch x_4, x_5, x_6, x_7 und y_4, y_5, y_6, y_7 bezeichnet werden, also der Anordnung

$$x_4 < x_2 < x_5 < x_1 < x_6 < x_3 < x_7 ; \quad y_4 < y_2 < y_5 < y_1 < y_6 < y_3 < y_7$$

genügen. So fahre man fort, so wird dadurch die Abbildung der beiden Mengen bewirkt. Es ist wiederum nur zu zeigen, daß in sie auch alle Punkte von R und M eingehen; die Erhaltung der Ordnung ist evident. Aber auch das erste ist leicht ersichtlich ... “

Die hier von Schoenflies geschilderte Konstruktion von Fréchet ist nicht identisch mit der Konstruktion im Beweis oben. Wesentlich gemeinsam ist ihnen aber, daß die beiden Ordnungen symmetrisch behandelt werden. Die einfachste symmetrische Behandlung ist vielleicht, daß nach n Schritten die ersten n Elemente der beiden Aufzählungen versorgt worden sind – wie im Beweis oben, nicht aber bei Schoenflies/Fréchet. Ein „Hin-und-Her“ ist aber nicht unbedingt notwendig für einen Beweis des Satzes von Cantor:

Übung

Seien wieder $x_0, x_1, \dots$ und $q_0, q_1, \dots$ Aufzählungen von M bzw. $\mathbb{Q}$ (ohne Wiederholungen). Wir definieren $f: M \rightarrow \mathbb{Q}$ durch Rekursion über $n \in \mathbb{N}$ wie folgt:

Sei $f(x_0) = q_0$. Ist $f(x_n)$ definiert, so setzen wir $f(x_{n+1}) = q_\ell$, wobei ℓ minimal gewählt wird, sodaß f weiterhin ordnungserhaltend ist.

Zeigen Sie, daß die so definierte Funktion $f : M \rightarrow \mathbb{Q}$ ein Ordnungsisomorphismus ist. [In dieser Weise hat Cantor den Satz bewiesen. Nichttrivial ist nur die Surjektivität von f . Die Annahme eines kleinsten ℓ mit $y_\ell \notin \text{rng}(f)$ führt zu einem Widerspruch.]

Cantor (1895): „Hat man eine einfach [linear] geordnete Menge M , welche die drei Bedingungen erfüllt:

- 1) [M ist abzählbar unendlich],
- 2) M hat kein dem Range nach niedrigstes und kein höchstes Element,
- 3) M ist überalldicht,

so ist der Ordnungstypus von M gleich $\eta \dots$ “

Diesen Charakterisierungssatz hat Cantor implizit bereits im sechsten Teil der „Linearen Punktmannigfaltigkeiten“ (Cantor, 1884b) bei der Untersuchung der Mächtigkeiten perfekter Mengen verwendet.

Für überabzählbare lineare Ordnungen ist ein zum Satz von Cantor analoger Charakterisierungssatz nicht mehr gültig:

Übung

Sei κ eine unendliche Kardinalzahl mit $\kappa \geq \omega_1$.

Dann existieren lineare Ordnungen $\langle M, < \rangle$ und $\langle N, < \rangle$ mit:

- (i) $|M| = |N| = \kappa$,
- (ii) $\langle M, < \rangle$ und $\langle N, < \rangle$ sind dicht und unbeschränkt,
- (iii) $\langle M, < \rangle$ und $\langle N, < \rangle$ sind nicht ordnungsisomorph.

Läßt man im Charakterisierungssatz die Forderung „unbeschränkt“ fallen, so sind vier Fälle möglich:

Korollar (Typen abzählbarer dichter Ordnungen)

Sei $\langle M, < \rangle$ eine abzählbare dichte lineare Ordnung, und sei $\alpha = \text{o.t.}(\langle M, < \rangle)$.

Dann gilt: $\alpha \in \{ \eta, 1 + \eta, \eta + 1, 1 + \eta + 1 \}$.

Beweis

Für eine beliebige lineare Ordnung $\langle N, < \rangle$ sei $\langle N', < \rangle$ die Ordnung, die aus N durch Entfernen evtl. vorhandener kleinster und größter Punkte entsteht. Ist $\text{o.t.}(\langle N', < \rangle) = \lambda$, so gilt offenbar

$\text{o.t.}(\langle N, < \rangle) \in \{ \lambda, 1 + \lambda, \lambda + 1, 1 + \lambda + 1 \}$.

Hieraus folgt die Behauptung, denn es gilt $\text{o.t.}(\langle N', < \rangle) = \eta$

– nach dem Charakterisierungssatz.

Wir stellen einige Rechenregeln für η zusammen. Besonders bemerkenswert ist (ii): Das Produkt von η mit einem beliebigen abzählbaren Typus α ist wieder η .

Übung

- (i) $\eta + \eta = \eta$,
- (ii) $\eta \cdot \alpha = \eta$ für alle abzählbaren Typen α ; insbesondere gilt $\eta \cdot \eta = \eta$,
- (iii) $\eta + 1 + \eta = \eta$.

Lücken in linearen Ordnungen

Weitere Anwendungen des Charakterisierungssatzes hängen mit dem Begriff der Lücke in einer linearen Ordnung zusammen. Die Ideen dieses Begriffs wurden zuerst von Richard Dedekind untersucht.

Die reellen Zahlen bilden anschaulich und begrifflich ein Kontinuum, sie haben keine Löcher. Dies besagt gerade die Vollständigkeitsbedingung, die wir allgemein für lineare Ordnungen definieren können.

Definition (*vollständige lineare Ordnungen*)

Sei $\langle M, < \rangle$ eine lineare Ordnung. Weiter sei $X \subseteq M$.

$s \in M$ heißt *obere Schranke von X* (bzgl. $<$), falls $X \leq s$.

$s \in M$ heißt *Supremum von X* (bzgl. $<$), in Zeichen $s = \sup(X)$, falls gilt: $X \leq s$, und für alle $s' \in M$ mit $X \leq s'$ gilt $s \leq s'$.

Analog sind untere Schranken und Infima $s = \inf(X)$ definiert.

$X \subseteq M$ heißt *nach oben beschränkt*, falls ein $s \in M$ existiert mit $X \leq s$.

Analog ist *nach unten beschränkt* definiert.

$\langle M, < \rangle$ heißt *vollständig*, falls jede nichtleere nach oben beschränkte Teilmenge X von M ein Supremum besitzt.

In einer vollständigen linearen Ordnung hat dann jede nach unten beschränkte nichtleere Teilmenge automatisch ein Infimum. Zur Diskussion der Vollständigkeit und ihrer Verletzung ist nun folgender Begriff nützlich:

Definition (*Dedekindscher Schnitt*)

Sei $\langle M, < \rangle$ eine lineare Ordnung.

Ein (*Dedekindscher*) *Schnitt* in $\langle M, < \rangle$ ist ein Paar (L, R) ,

$L, R \subseteq M$ mit den Eigenschaften:

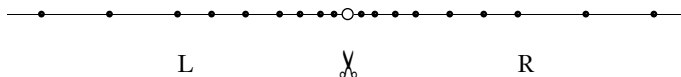
- (i) $L, R \neq \emptyset, L \cap R = \emptyset, L \cup R = M,$
 - (ii) für alle $x \in L, y \in R$ gilt $x < y,$
 - (iii) $\sup(L) \in L$, falls $\sup(L)$ existiert.
- $\langle M, < \rangle$

L
 $\nexists$
 R

Ein Schnitt ist also eine Zerlegung der „Linie“ M in einen linken Teil L und einen rechten Teil R . Per Konvention enthält der linke Teil sein Supremum, falls dieses existiert. Falls dieses Supremum – die „Schnittstelle“ zwischen L und R – nicht existiert, so sprechen wir von einer Lücke der Ordnung:

Definition (*Lücken in linearen Ordnungen*)

Sei $\langle M, < \rangle$ eine lineare Ordnung, und sei (L, R) ein Schnitt in $\langle M, < \rangle$.
 (L, R) heißt eine *Lücke* in $\langle M, < \rangle$, falls $\sup(L)$ nicht existiert.



Nach dieser Definition hat die Ordnung $\langle \mathbb{N}, < \rangle$ keine Lücken. Eine Lücke meint eher eine „Punktierung“. Sprungstellen zwischen einem Element und seinem direkten Nachfolger gelten nicht als Lücken:

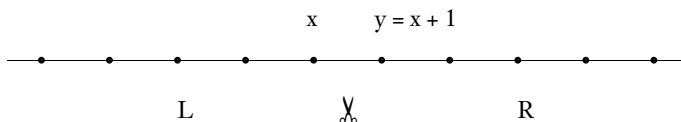
Definition (*Sprungstelle einer linearen Ordnung*)

Sei $\langle M, < \rangle$ eine lineare Ordnung, und sei (L, R) ein Schnitt in $\langle M, < \rangle$.
 (L, R) heißt eine *Sprungstelle* in $\langle M, < \rangle$, falls gilt:

$x = \sup(L)$ und $y = \inf(R)$ existieren und es gilt $x < y$.

y heißt dann der *Nachfolger von x* in $\langle M, < \rangle$, in Zeichen $y = x + 1$.

Analog heißt x der *Vorgänger von y* in $\langle M, < \rangle$, in Zeichen $x = y - 1$.



Eine lineare Ordnung ist offenbar genau dann dicht, wenn sie keine Sprungstellen besitzt, und sie ist genau dann vollständig, wenn sie keine Lücken aufweist. Die Sprünge und Lücken einer linearen Ordnung lokalisieren genau die Stellen, an denen die Dichtheit bzw. die Vollständigkeit verletzt ist.

Hausdorff (1914): „Eine Zerlegung [des Trägers A einer linearen Ordnung $\langle A, < \rangle$] in zwei nicht verschwindende Stücke

$$A = P + Q \quad [P, Q \neq \emptyset, A = P \cup Q, P \cap Q = \emptyset, x < y \text{ für alle } x \in P, y \in Q]$$

möge ein Sprung heißen, wenn P ein letztes Element und Q ein erstes hat; ein Schnitt, wenn P ein letztes und Q kein erstes, oder umgekehrt P kein letztes und Q ein erstes Element hat; eine Lücke, wenn weder P ein letztes noch Q ein erstes Element hat.“

Diese Definitionen stimmen mit den unseren überein, wobei wir auch Sprünge und Lücken als Schnitte gelten lassen, und bei einem Schnitt (L, R) der Bequemlichkeit halber immer $\sup(L) \in L$ fordern, falls das Supremum existiert.

Übung

Seien $\langle M, < \rangle$ und $\langle N, < \rangle$ lineare Ordnungen.

Bestimmen Sie alle Lücken und Sprungstellen von $\langle M, < \rangle \cdot \langle N, < \rangle$.

Eine wichtige Beobachtung ist nun, daß ähnliche Ordnungen ähnliche Lücken haben:

Satz (*Lücken ähnlicher Ordnungen*)

Seien $\langle M, < \rangle$ und $\langle N, < \rangle$ ähnliche lineare Ordnungen, und sei $f: M \rightarrow N$ ein Ordnungsisomorphismus.

- (i) Ist $X \subseteq M$ und ist $x = \sup(X)$ in M , so ist $f(x) = \sup(f''X)$ in N .
Analog folgt aus $x = \inf(X)$ in M , daß $f(x) = \inf(f''X)$ in N .
- (ii) Ist (L, R) eine Lücke in $\langle M, < \rangle$, und sind $L' = f''L$, $R' = f''R$, so ist (L', R') eine Lücke in $\langle N, < \rangle$.
- (iii) $\langle M, < \rangle$ hat eine Lücke *gdw* $\langle N, < \rangle$ hat eine Lücke.

Beweis

zu (i): Sei $X' = f''X$. Wegen f ordnungserhaltend gilt $X' \leq f(x)$.

Annahme, es existiert ein $y' \in N$ mit $X' \leq y'$ und $y' < f(x)$.

Sei $y = f^{-1}(y')$. Wegen $y' = f(y) < f(x)$ ist dann $y < x$.

Wegen $x = \sup(X)$ existiert also ein $z \in X$ mit $y < z \leq x$.

Dann ist aber $y' = f(y) < f(z) \in X'$, im Widerspruch zu $X' \leq y'$.

zu (ii): Offenbar ist (L', R') ein Schnitt in $\langle N, < \rangle$.

Annahme, (L', R') ist keine Lücke in $\langle N, < \rangle$.

Dann existiert $x' = \sup(L')$ in $\langle N, < \rangle$.

Da $f^{-1}: N \rightarrow M$ ein Ordnungsisomorphismus ist, existiert nach (i) $\sup(f^{-1}''L') \in L$, im Widerspruch zu (L, R) Lücke in $\langle M, < \rangle$.

– zu (iii): Folgt aus (i) und (ii).

Ein Ordnungsisomorphismus übersetzt also einen Schnitt in einen Schnitt und erhält dabei Lücken. Ist von zwei ähnlichen Ordnungen die eine vollständig, so ist auch die andere vollständig.

Die Existenz von Lücken in $\langle \mathbb{Q}, < \rangle$ folgt aus der Existenz irrationaler Zahlen. Sei etwa $w \in \mathbb{R}$ die positive Quadratwurzel aus 2, und seien $L = \{x \in \mathbb{Q} \mid x < w\}$, $R = \{x \in \mathbb{Q} \mid w < x\}$. Dann ist (L, R) eine Lücke von $\mathbb{Q}$, denn w ist irrational. Analog folgt die Existenz von Lücken für die algebraischen Zahlen $\langle \mathbb{A}, < \rangle$ aus der Existenz von transzendenten Zahlen.

Aus dem Satz von Cantor über die rationalen Zahlen folgt nun aber rein ordnungstheoretisch – ohne irgendwelche Zahlentheorie! – die Existenz von Lücken in $\langle \mathbb{Q}, < \rangle$ und allen abzählbaren Erweiterungen der rationalen Zahlen:

Satz (*Existenz von Lücken in $\mathbb{Q}$*)

- (i) Sei $\langle M, < \rangle$ eine lineare Ordnung mit $\eta = \text{o.t.}(\langle M, < \rangle)$.
Dann hat $\langle M, < \rangle$ Lücken.
- (ii) Ist $A \subseteq \mathbb{R}$ abzählbar und gilt $\mathbb{Q} \subseteq A$, so hat $\langle A, < \rangle$ Lücken (unter der natürlichen Ordnung von $\mathbb{R}$).

Beweis

- zu (i): Da ähnliche Ordnungen das gleiche Lückenverhalten haben, genügt es, *eine* lineare Ordnung vom Typ η mit Lücken anzugeben. Sei $\mathbb{Q}' = \mathbb{Q} - \{0\}$, und seien $L = \{q \in \mathbb{Q} \mid q < 0\}$, $R = \{q \in \mathbb{Q} \mid 0 < q\}$. Dann ist (L, R) eine Lücke in $\langle \mathbb{Q}', < \rangle$. Aber $\langle \mathbb{Q}', < \rangle$ ist abzählbar, unbeschränkt und dicht, nach dem Charakterisierungssatz gilt also $\eta = \text{o. t.}(\langle \mathbb{Q}', < \rangle)$. Dies zeigt die Behauptung.
- zu (ii): Nach Voraussetzung ist A abzählbar, und wegen $\mathbb{Q} \subseteq A \subseteq \mathbb{R}$ ist $\langle A, < \rangle$ unbeschränkt und dicht. Also gilt $\eta = \text{o. t.}(\langle A, < \rangle)$,
 — und die Behauptung folgt damit aus (i).

Der Beweis von (i) lautet im wesentlichen: $\eta + \eta = \eta$, und $\eta + \eta$ hat offenbar eine Lücke. Die algebraischen Zahlen haben Lücken nach (ii). Diese Lücken sind genau die transzendenten Zahlen.

Korollar (*Überabzählbarkeit von $\mathbb{R}$*)

Die reellen Zahlen sind überabzählbar.

Beweis

- Wäre $\mathbb{R}$ abzählbar, so hätte $\langle \mathbb{R}, < \rangle$ nach (ii) im Satz oben (mit $A = \mathbb{R}$)
 — Lücken. $\langle \mathbb{R}, < \rangle$ ist aber vollständig. Also ist $\mathbb{R}$ überabzählbar.

Als Ganzes verläuft dieser Beweis der Überabzählbarkeit von $\mathbb{R}$ wie folgt:

1. $\langle \mathbb{Q}, < \rangle$ ist unvollständig. Denn $\langle \mathbb{Q} - \{0\}, < \rangle$ ist unvollständig und ähnlich zu $\langle \mathbb{Q}, < \rangle$, und ähnliche Ordnungen haben ähnliche Lücken.
2. Wäre $\mathbb{R}$ abzählbar, so wäre $\langle \mathbb{R}, < \rangle$ ähnlich zu $\langle \mathbb{Q}, < \rangle$, also ebenfalls unvollständig, *Widerspruch!* Also ist $\mathbb{R}$ überabzählbar.

Hierbei wird zweimal der Charakterisierungssatz über den Typ η benutzt: Für die Unvollständigkeit von $\langle \mathbb{Q}, < \rangle$ und für „ $\mathbb{R}$ abzählbar folgt $\langle \mathbb{R}, < \rangle = \langle \mathbb{Q}, < \rangle$ “.

Die Unvollständigkeit von $\mathbb{Q}$ kann man natürlich direkt benutzen, wenn man die Existenz irrationaler Zahlen als bekannt voraussetzt. Der ordnungstheoretische Existenzbeweis irrationaler Zahlen via $\langle \mathbb{Q} - \{0\}, < \rangle \equiv \langle \mathbb{Q}, < \rangle$ ist aber für sich von Interesse, und zeigt alleine schon die Stärke des Charakterisierungssatzes. Die Existenz von Lücken ist dem Ordnungstyp η in die Wiege gelegt. Und die Lücken von $\mathbb{Q}$ sind die irrationalen Zahlen... Ein Beweis ganz ohne Arithmetik, aber ganz in der Tradition des Irritierenden, das den Existenzbeweisen irrationaler Zahlen seit jeher anhaftet.

Es ist überraschend, daß sich dieses letztendlich einfache Korollar zu Cantors Charakterisierungssatz der rationalen Zahlen nur sehr selten in der Literatur findet. Bei Hausdorff liest man:

Hausdorff (1914): „Abzählbare stetige Typen [Typen vollständiger und dichter Ordnungen] gibt es überhaupt nicht; denn stetige Typen sind ja dicht, und die vier abzählbaren dichten Typen $[\eta, 1 + \eta, \eta + 1, 1 + \eta + 1]$ sind unstetig.“

Wir haben also einen rein ordnungstheoretischen Beweis der Überabzählbarkeit der reellen Zahlen erhalten. Genauer haben wir gezeigt: Ein *Kontinuum* $\langle \mathbb{K}, < \rangle$ ist notwendig überabzählbar, wobei wir hier an den Begriff des Kontinuums die folgenden Anforderungen stellen: 1. Ein Kontinuum ist dicht (und hat mindestens zwei Elemente). 2. Ein Kontinuum hat keine Lücken.

Einbettungen

Wir zeigen, daß im Reich der abzählbaren Ordnungstypen der Typ η der rationalen Zahlen das Maß aller Dinge ist. Hierzu ein natürlicher Begriff.

Definition (Einbettung)

Seien $\langle M, < \rangle$ und $\langle N, < \rangle$ lineare Ordnungen.

- (i) $f : M \rightarrow N$ heißt eine *Einbettung* von $\langle M, < \rangle$ in $\langle N, < \rangle$, falls für alle $x, y \in M$ gilt:
 $x < y$ gdw $f(x) < f(y)$.

f heißt *korrekt*, falls zusätzlich für alle $X \subseteq M$ gilt:

- (a) Ist $x = \sup(X)$ in M , so ist $f(x) = \sup(f''X)$ in N .
 (b) Ist $x = \inf(X)$ in M , so ist $f(x) = \inf(f''X)$ in N .

- (ii) $\langle M, < \rangle$ läßt sich in $\langle N, < \rangle$ (*korrekt*) *einbetten*, falls eine (korrekte) Einbettung f von $\langle M, < \rangle$ in $\langle N, < \rangle$ existiert.

Ist $f : M \rightarrow N$ eine Einbettung von $\langle M, < \rangle$ in $\langle N, < \rangle$ mit $\text{rng}(f) = N'$, so ist $f : M \rightarrow N'$ ein Ordnungsisomorphismus von $\langle M, < \rangle$ nach $\langle N', < \rangle$. Dieser Ordnungsisomorphismus erhält Suprema und Infima, aber Suprema in $\langle N', < \rangle$ fallen i. a. nicht mit Suprema in $\langle N, < \rangle$ zusammen. Für korrekte Einbettungen ist dies aber der Fall.

Ist z. B. $N = \mathbb{R}$, $A = \{ -1/n \mid n \in \mathbb{N}, n \geq 1 \}$ und $N' = A \cup \{ 1 \}$, so ist $\sup(A) = 1$ in $\langle N', < \rangle$, aber $\sup(A) = 0$ in $\langle N, < \rangle$.

Definition ($\alpha \leq \beta$ und $\alpha \leq^* \beta$)

Seien α, β Ordnungstypen. Wir setzen:

- $\alpha \leq \beta$, falls eine Einbettung f von $\langle M, < \rangle$ in $\langle N, < \rangle$ existiert, wobei $\langle M, < \rangle, \langle N, < \rangle$ lineare Ordnungen sind mit o. t. $(\langle M, < \rangle) = \alpha$, o. t. $(\langle N, < \rangle) = \beta$.

- $\alpha \leq^* \beta$, falls eine korrekte derartige Einbettung f existiert.

Übung

- (i) $\leq$ und $\leq^*$ sind reflexiv und transitiv.
 (ii) Aus $\alpha \leq^* \beta$ und $\beta \leq^* \alpha$ folgt i. a. nicht $\alpha = \beta$.
 (iii) Es gibt α, β mit $\alpha \leq \beta$ und $\text{non}(\alpha \leq^* \beta)$.

Aus dem Charakterisierungssatz erhalten wir nun, daß der Typus η ein Dach für alle abzählbaren Ordnungstypen darstellt:

Satz (*Universalität des Typs η*)

Sei α ein abzählbarer Ordnungstyp.

Dann gilt $\alpha \leq^* \eta$.

$$\eta \begin{array}{c} * \\ \swarrow \end{array} \quad * \quad * \searrow$$

abzählbare Typen

Beweis

Sei $\langle M, < \rangle$ eine lineare Ordnung des Typs α .

Weiter sei $\langle N, < \rangle = \langle \mathbb{Q}, < \rangle + \langle M, < \rangle + \langle \mathbb{Q}, < \rangle$.

Dann ist $\langle N, < \rangle$ abzählbar und unbeschränkt. Wir erweitern $\langle N, < \rangle$ zu einer dichten Ordnung $\langle Q, <_Q \rangle$, indem wir an allen Sprungstellen der Ordnung eine Kopie von $\mathbb{Q}$ einschieben:

Sei $S = \{x \in N \mid x+1 \text{ existiert in } N\}$.

Wir setzen

$$Q = N \cup (S \times \mathbb{Q}),$$

$$\{x\} \times \mathbb{Q}$$

wobei o. E. $N \cap (S \times \mathbb{Q}) = \emptyset$.

$$x \searrow y$$

Die Ordnung $<_Q$ ist definiert durch:

$$y = x + 1 \text{ in } \langle N, < \rangle$$

$$(i) \quad <_N \subseteq <_Q,$$

$$(ii) \quad (x, q_1) <_Q (y, q_2), \quad \text{falls} \quad \begin{array}{l} x <_N y \text{ oder} \\ x = y \text{ und } q_1 <_Q q_2, \end{array}$$

$$(iii) \quad (x, q) <_Q y, \quad \text{falls} \quad x <_N y,$$

$$(iv) \quad x <_Q (y, q), \quad \text{falls} \quad x \leq_N y.$$

Dann gilt o. t. $\langle Q, <_Q \rangle = \eta$. Also existiert ein Ordnungsisomorphismus

$$g: Q \rightarrow \mathbb{Q}.$$

Dann ist aber $f = g|_M$ eine korrekte Einbettung von $\langle M, < \rangle$ in $\langle \mathbb{Q}, < \rangle$:

Offenbar ist f eine Einbettung. Ist nun $X \subseteq M$ und existiert $x = \sup(X)$ in M , so ist nach Konstruktion von $\langle Q, <_Q \rangle$ auch $x = \sup(X)$ in Q , und es gilt $g(x) = \sup(g''X)$, da g ein Ordnungsisomorphismus ist. Also auch $f(x) = \sup(f''X)$ wegen $f = g|_M$. Analoges gilt für Infima.

— Also ist f korrekt, und damit gilt $\alpha \leq^* \eta$.

$\langle \mathbb{Q}, < \rangle$ – und allgemein jede lineare Ordnung des Typs η – enthält also eine korrekte Kopie jeder abzählbaren linearen Ordnung. Insbesondere existiert für jede abzählbare Ordinalzahl α eine strikt aufsteigende Folge rationaler Zahlen der Länge α :

Korollar (*lange aufsteigende Folgen in $\mathbb{Q}$*)

Sei α eine abzählbare Ordinalzahl.

Dann existiert eine strikt aufsteigende stetige Folge $\langle q_\beta \mid \beta < \alpha \rangle$ rationaler Zahlen, d. h. es gilt:

$$(i) \quad \beta < \gamma \text{ gdw } q_\beta < q_\gamma \quad \text{für alle } \beta, \gamma < \alpha,$$

$$(ii) \quad q_\lambda = \sup(\{q_\beta \mid \beta < \lambda\}) \quad \text{für Limesordinalzahlen } \lambda < \alpha.$$

Beweis

$\langle W(\alpha), < \rangle$ ist eine abzählbare lineare Ordnung.

Also existiert eine korrekte Einbettung $f : W(\alpha) \rightarrow \mathbb{Q}$.

- Dann ist $f = \langle q_\beta \mid \beta < \alpha \rangle$ wie gewünscht.

Man kann also alle abzählbaren Ordinalzahlen durch Teilordnungen von $\mathbb{Q}$ visualisieren. Die reellen Zahlen leisten hier nicht mehr als die rationalen Zahlen. Auch wenn wir sie zugrunde legen, ist eine Visualisierung durch Einbettung für überabzählbare Ordinalzahlen nicht mehr möglich: Es gibt keine strikt aufsteigenden Folgen der Länge ω_1 in $\mathbb{R}$. Denn ist $\langle r_\beta \mid \beta < \alpha \rangle$ strikt aufsteigend in $\mathbb{R}$, so ist $\mathbb{Q} \cap]r_\beta, r_{\beta+1}[\neq \emptyset$ für alle β mit $\beta + 1 < \alpha$. Wegen der Abzählbarkeit von $\mathbb{Q}$ ist also α notwendig abzählbar.

Weiter erhalten wir auch für jeden abzählbaren Ordnungstyp α die Existenz einer transzendenten Teilmenge von $\mathbb{R}$ des Typs α , und wir können auch hier wieder eine korrekte Einbettung erreichen:

Korollar (*transzendente Teilmengen von $\mathbb{R}$*)

Sei $\langle M, < \rangle$ eine abzählbare lineare Ordnung.

Dann existiert ein $f : M \rightarrow \mathbb{R}$ mit:

- (i) f ist eine korrekte Einbettung von $\langle M, < \rangle$ in $\langle \mathbb{R}, < \rangle$,
- (ii) $f(x)$ ist transzendent für alle $x \in M$.

Beweis

Für $n \in \mathbb{N}$, $n \neq 0$, und $k \in \mathbb{Z}$ sei

$x_{n,k} =$ „eine transzendente Zahl z mit $z \in [k/n, (k+1)/n]$ “,

und es sei dann

$T = \{ x_{n,k} \mid n \in \mathbb{N} - \{0\}, k \in \mathbb{Z} \}$.

Dann ist T eine Menge von transzendenten Zahlen mit o. t. $\langle T, < \rangle = \eta$.

Nach dem Satz oben existiert dann eine korrekte Einbettung $f : M \rightarrow T$

von $\langle M, < \rangle$ in $\langle T, < \rangle$. T ist aber dicht in $\mathbb{R}$, und damit gilt für alle $X \subseteq T$:

Ist $x = \sup(X)$ in $\langle T, < \rangle$, so ist $x = \sup(X)$ in $\langle \mathbb{R}, < \rangle$.

Analoges gilt für Infima.

- Also ist f auch eine korrekte Einbettung von $\langle M, < \rangle$ in $\langle \mathbb{R}, < \rangle$.

Insbesondere existiert für jede abzählbare Ordinalzahl α eine Menge T von transzendenten Zahlen mit o. t. $\langle T, < \rangle = \alpha + 1$ und $\sup(X) \in T$ für alle nichtleeren Teilmengen X von T .

Mit dieser Untersuchung von η sind wir nun bestens gerüstet für eine ordnungstheoretische Charakterisierung der reellen Zahlen.

Der Ordnungstyp θ

Wir charakterisieren nun den Ordnungstyp der reellen Zahlen. Auch dieses Ergebnis geht auf Cantor zurück, der die folgende bemerkenswerte Eigenschaft der reellen Zahlen unter dem Scheffel hervorgerückt hat: Obwohl die reellen Zahlen überabzählbar sind, besitzen sie in den rationalen Zahlen eine abzählbare dichte Teilmenge – zwischen zwei verschiedenen reellen Zahlen liegt immer eine rationale Zahl. Für diese Eigenschaft, die man allgemein für lineare Ordnungen formulieren kann, hat sich ein eigener Begriff durchgesetzt:

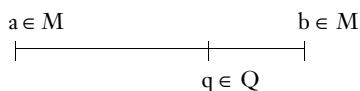
Definition (*dichte Teilmengen und separable lineare Ordnungen*)

Sei $\langle M, < \rangle$ eine lineare Ordnung.

$Q \subseteq M$ heißt *dicht in* $\langle M, < \rangle$, falls gilt:

Für alle $a, b \in M$ mit $a < b$

existiert ein $q \in Q$ mit $a < q < b$.



$\langle M, < \rangle$ heißt *separabel*, falls eine abzählbare dichte Teilmenge von M existiert.

Für die reellen Zahlen gilt:

$\langle \mathbb{R}, < \rangle$ ist unbeschränkt, vollständig und separabel.

Diese drei Bedingungen charakterisieren bereits die Ordnung $\langle \mathbb{R}, < \rangle$:

Satz (*Charakterisierung von θ*)

Sei $\langle M, < \rangle$ eine lineare Ordnung. Es gelte:

- (i) $\langle M, < \rangle$ ist unbeschränkt,
- (ii) $\langle M, < \rangle$ ist vollständig,
- (iii) $\langle M, < \rangle$ ist separabel.

Dann gilt $\theta = o.t.(\langle M, < \rangle)$, d. h. $\langle M, < \rangle$ und $\langle \mathbb{R}, < \rangle$ sind ähnlich.

Beweis

Sei Q eine abzählbare dichte Teilmenge von M . Dann ist Q unbeschränkt. Nach dem Satz von Cantor über den Ordnungstyp η sind also $\langle Q, < \rangle$ und $\langle \mathbb{Q}, < \rangle$ ähnlich.

Sei $f : Q \rightarrow \mathbb{Q}$ ein Ordnungsisomorphismus. Wir setzen f zu einer Funktion $g : M \rightarrow \mathbb{R}$ fort. Da g ein Ordnungsisomorphismus werden soll, müssen Suprema von Teilmengen von M auf die Suprema der Bildmengen abgebildet werden. Da Q in M dicht ist, ist jedes Element von M Supremum einer Teilmenge von Q . Wir definieren also für $x \in M$:

$$g(x) = \sup(\{f(q) \mid q \in Q, q < x\}).$$

Wegen Q dicht gilt für alle $x \in Q$, daß $x = \sup(\{q \mid q \in Q, q < x\})$.

Da $f : Q \rightarrow \mathbb{Q}$ ein Ordnungsisomorphismus ist, gilt also

$f(x) = \sup(\{f(q) \mid q \in \mathbb{Q}, q < x\})$ für alle $x \in \mathbb{Q}$.

Somit ist $g(x) = f(x)$ für alle $x \in \mathbb{Q}$, d. h. g ist eine Fortsetzung von f .

g ist ordnungserhaltend und damit injektiv:

Seien $x, y \in M$, $x < y$.

Sei $z \in \mathbb{Q}$ mit $x < z < y$. Dann ist $g(x) < f(z) < g(y)$.

g ist aber auch surjektiv:

Denn sei $x' \in \mathbb{R}$ und $X' = \{q \in \mathbb{Q} \mid q < x'\}$.

Dann ist $x' = \sup(X')$. Sei $X = f^{-1}X' \subseteq \mathbb{Q}$ und $x = \sup(X)$.

Dann ist $g(x) = \sup(\{f(q) \mid q \in X\}) = \sup(\{q \mid q \in X'\}) = x'$.

Also ist $g : M \rightarrow \mathbb{R}$ ein Ordnungsisomorphismus.

– $\langle M, < \rangle$ und $\langle \mathbb{R}, < \rangle$ sind also ähnlich, und somit gilt $\theta = o. t.(\langle M, < \rangle)$.

Lassen wir die Unbeschränktheit fallen, so sind vier Fälle möglich:

Übung

Sei $\langle M, < \rangle$ eine vollständige separable lineare Ordnung mit mehr als einem Element.

Dann ist $\langle M, < \rangle$ ordnungsisomorph zu $\langle I, < \rangle$, wobei

$I = [0, 1]$ oder $I = [0, 1[$ oder $I =]0, 1]$ oder $I =]0, 1[$.

[Unterscheide nach der Existenz von kleinsten und größten Elementen in $\langle M, < \rangle$. $\langle \mathbb{R}, < \rangle$ ist ähnlich zu $\langle]0, 1[, < \rangle$.]

Aus dem Beweis erhalten wir einen Einbettungssatz:

Korollar (Einbettungssatz für $\mathbb{R}$)

Sei $\langle M, < \rangle$ eine separable lineare Ordnung.

Dann läßt sich $\langle M, < \rangle$ in $\langle \mathbb{R}, < \rangle$ korrekt einbetten.

Es gilt also $\alpha \leq^* \theta$ für alle separablen Typen α .

Cantor hat die Charakterisierung der Ordnung der reellen Zahlen etwas anders formuliert. Er betrachtet $X = [0, 1] \subseteq \mathbb{R}$ und charakterisiert $\langle X, < \rangle$ wie folgt:

Cantor (1895): „Hat eine geordnete Menge M ein solches Gepräge, daß sie 1) ‚perfekt‘ ist, 2) in ihr eine Menge S mit der Kardinalzahl ... $\aleph_0$ enthalten ist, welche zu M in der Beziehung steht, daß zwischen je zwei beliebigen Elementen m_0 und m_1 von M Elemente von S dem Range nach liegen, so ist $[o. t.(\langle M, < \rangle) = o. t.(\langle [0, 1], < \rangle)]$.“

Die Eigenschaft *perfekt* bedeutet hier bei Cantor: Jede aufsteigende bzw. absteigende abzählbare Folge in M hat ein Supremum bzw. ein Infimum in M , und jedes $x \in M$ ist Limes einer abzählbaren auf- oder absteigenden Folge in M .

Zermelo schreibt in einer Anmerkung zur Cantorsche Charakterisierung:

Zermelo (1932, Anmerkung zu § 11 von Cantor 1895): „Von den beiden Eigenschaften, durch welche hier Cantor den Ordnungstyp des Linearkontinuums charakterisiert, ist von grundlegender Bedeutung und die eigentliche Cantorsche Entdeckung die von ihm als 2) bezeichnete Eigenschaft: die Existenz einer in M ‚überall dichten‘ abzählbaren Teilmenge S . Weniger glücklich erscheint uns heute seine Formulierung der Eigenschaft 1) ...“

Zermelo schlägt statt 1) die Vollständigkeitsbedingung „ $\langle M, < \rangle$ hat keine Lücken“ vor, die unserer Formulierung entspricht.

Dedekind-Vervollständigung einer linearen Ordnung

Die Diskussion der Lücken einer linearen Ordnung und der Charakterisierungssatz legen die folgende Konstruktion einer Ordnung des Typs θ nahe (Dedekind, 1872): Wir schließen alle Lücken einer Ordnung des Typs η . Ein solches Verfahren läßt sich für jede lineare Ordnung durchführen.

Definition (*Dedekind-Vervollständigung einer linearen Ordnung*)

Sei $\langle M, < \rangle$ eine unbeschränkte lineare Ordnung.

Sei $\mathcal{D}(M) = \{ (L, R) \mid (L, R) \text{ ist ein Dedekindscher Schnitt von } M \}$.

Für $(L, R), (L', R') \in \mathcal{D}(M)$ setzen wir

$(L, R) < (L', R') \text{ falls } L \subset L'.$

Dann heißt $\langle \mathcal{D}(M), < \rangle$ die (*Dedekind-*)*Vervollständigung* von $\langle M, < \rangle$.

Wir nehmen hier der Bequemlichkeit halber „unbeschränkt“ an, denn nach unserer Schnittdefinition ist $(M, \emptyset)$ kein Schnitt, da der zweite Teil leer ist. Läßt man $(M, \emptyset)$ oder $(\emptyset, M)$ als Schnitt zu, so erhält man bei der Vervollständigung einen zusätzlichen rechten oder linken Randpunkt, falls M nach rechts oder links unbeschränkt ist.

Ist $x \in M$ und (L_x, R_x) der Schnitt mit $L_x = \{ y \in M \mid y \leq x \}$, so kann man x mit (L_x, R_x) identifizieren, und damit $M \subseteq \mathcal{D}(M)$ annehmen. Mit dieser Einbettung können wir die Dedekind-Vervollständigung einer linearen Ordnung so beschreiben: Wir schließen alle Lücken der Ordnung, und als lückenfüllende Objekte verwenden wir gerade die Schnitte der Ordnung selbst, welche eine Lücke der Ordnung bestimmen.

Satz (*Vollständigkeit der Dedekind-Vervollständigung*)

Sei $\langle M, < \rangle$ eine unbeschränkte lineare Ordnung, und sei $\langle \mathcal{D}(M), < \rangle$ die Dedekind-Vervollständigung von $\langle M, < \rangle$.

Dann ist $\langle \mathcal{D}(M), < \rangle$ eine unbeschränkte und vollständige lineare Ordnung.

Ist $\langle M, < \rangle$ separabel, so gilt o. t. $\langle \mathcal{D}(M), < \rangle = \theta$.

Beweis

$\langle \mathcal{D}(M), < \rangle$ ist offenbar eine unbeschränkte lineare Ordnung.

Sei also $X \subseteq \mathcal{D}(M)$ nichtleer und beschränkt in $\langle \mathcal{D}(M), < \rangle$. Wir zeigen, daß X ein Infimum besitzt. Hierzu setzen wir:

$$L^* = \bigcap \{ L \mid (L, R) \in X \}, \quad R^* = M - L^*.$$

Dann ist (L^*, R^*) ein Schnitt in $\langle M, < \rangle$, und es gilt $(L^*, R^*) = \inf(X)$, denn für alle $(L', R') \in \mathcal{D}(M)$ gilt:

$$(L', R') \leq X \text{ gdw } L' \subseteq L \text{ für alle } (L, R) \in X \text{ gdw } L' \subseteq L^*.$$

Also ist $\langle \mathcal{D}(M), < \rangle$ vollständig.

Sei nun $\langle M, < \rangle$ separabel, und sei $Q \subseteq M$ eine abzählbare dichte Teilmenge von $\langle M, < \rangle$. Wir nehmen $M \subseteq \mathcal{D}(M)$ an. Dann ist Q dicht in $\langle \mathcal{D}(M), < \rangle$ (!).

Also ist $\langle \mathcal{D}(M), < \rangle$ unbeschränkt, vollständig und separabel, und somit gilt

– o. t. $(\langle \mathcal{D}(M), < \rangle) = \theta$ nach dem Charakterisierungssatz für den Typ θ .

Insbesondere ist die Dedekind-Vervollständigung jeder linearen Ordnung des Typs η vom Typ θ und hat daher stets die Mächtigkeit von $\mathbb{R}$.

Nach dem Satz sind $\langle \mathbb{R}, < \rangle$ und $\langle \mathcal{D}(\mathbb{Q}), < \rangle$ ähnlich. In einem Aufbau des Zahlensystems kann man die reellen Zahlen und ihre Ordnung geradezu als die Vervollständigung von $\langle \mathbb{Q}, < \rangle$ definieren. Daß die vertrauten reellen Zahlen dann zu komplexen Schnitten $(L, R) \in \mathcal{P}(\mathbb{Q})^2$ werden, ist etwas schwer verdaulich. Man kann den Übergang von $\mathbb{Q}$ nach $\mathbb{R}$ aber sehr anschaulich zusammenfassen: wir füllen lediglich die Lücken von $\mathbb{Q}$ auf, mit was auch immer.

Hausdorff (1914): „Wir können diesem Verfahren [der Dedekind-Vervollständigung $\mathfrak{B} = \langle \mathcal{D}(A), < \rangle$ einer linearen Ordnung $\mathfrak{A} = \langle A, < \rangle$], freilich mit Preisgabe der eindeutigen Bestimmtheit, eine etwas anschaulichere Fassung geben: wir verschaffen uns eine zu A fremde Menge B , die mit der Menge $\mathfrak{B}$ [entspricht $\mathcal{D}(A) - A$] der lückenbestimmenden Anfangsstücke P [von $\mathfrak{A}$] äquivalent [gleichmächtig] ist, und übertragen die Ordnung von $\mathfrak{B} = \mathfrak{A} + \mathfrak{B}^1$ auf die äquivalente Menge $A + B$ [$A \cup B$], wobei die eineindeutige Beziehung zwischen A und $\mathfrak{A}$ durch die Zusammengehörigkeit von a und A^a [$\{ a' \in A \mid a' \leq a \}$] gegeben ist. Die Menge $A + B$ entsteht dann aus A durch Ausfüllung der Lücken, d. h. indem man in jede Lücke $P + Q$ ein Element b einschiebt ($P < b < Q$), wodurch die b gegenüber den a und auch untereinander geordnet werden.

Ist A die Menge der rationalen Zahlen, so ist B die Menge der irrationalen, $A + B$ die der reellen Zahlen. Was die Elemente b eigentlich ‚sind‘, bleibt dabei freilich unentschieden, ist aber auch gleichgültig, da es nur auf ihre Ordnung zu den a und untereinander (und auf ihre Rechengesetze) ankommt; übrigens kann man auch die Anfangsstücke P selber ‚reelle Zahlen‘ nennen, die schnittbestimmenden [solche mit einem Supremum] rationale reelle Zahlen, die lückenbestimmenden irrationalen reelle Zahlen, und dann, auf die Gefahr einer Verwechslung zwischen a und A^a hin, die Weglassung des Beiworts ‚reell‘ verabreden.

1 Die Summen sind hier nur im Sinne der [Vereinigung] ungeordneter Mengen zu verstehen; die Elemente von $\mathfrak{A}$ und $\mathfrak{B}$ liegen ja durcheinander [d. h. $\mathfrak{A} + \mathfrak{B}$ entsteht hier nicht etwa dadurch, daß $\mathfrak{A}$ um $\mathfrak{B}$ enderweitert wird.]“

Eine weitere Konstruktion von $\langle \mathbb{R}, < \rangle$ aus $\langle \mathbb{Q}, < \rangle$

Eine weitere, in der Analysis bevorzugte Konstruktionsmethode der reellen Zahlen benutzt konvergente Folgen rationaler Zahlen. Diese Methode hat gegenüber der ordnungstheoretischen Konstruktion den Vorteil, daß die Arithmetik auf $\mathbb{R}$ leichter aus der Arithmetik auf $\mathbb{Q}$ gewonnen werden kann.

Die Idee ist, eine reelle Zahl als eine konvergente Folge rationaler Zahlen aufzufassen. Zur Umsetzung dieser Idee ist es nötig, die Konvergenz einer Folge zu definieren, ohne ihren Limes bereits zu benutzen. Entscheidend ist hier der Begriff der Fundamentalfolge, den wir für die reellen Zahlen bereits kennengelernt haben.

Definition (Fundamentalfolge in $\mathbb{Q}$)

Sei $x = \langle x_n \mid n \in \mathbb{N} \rangle$ eine Folge rationaler Zahlen.

x heißt *Fundamentalfolge in $\mathbb{Q}$* , falls gilt:

Für alle $\varepsilon \in \mathbb{Q}$, $\varepsilon > 0$, existiert ein $n_0 \in \mathbb{N}$ mit:

$$|x_m - x_n| < \varepsilon \quad \text{für alle } n, m \in \mathbb{N} \text{ mit } n, m \geq n_0.$$

Zwei Fundamentalfolgen können den gleichen Grenzwert haben. Auch dies läßt sich ohne Vorgriff auf diesen Grenzwert definieren:

Definition (Nullfolge und „gleicher Limes“)

- (i) Eine Fundamentalfolge $\langle x_n \mid n \in \mathbb{N} \rangle$ heißt *Nullfolge*, falls gilt:

Für alle $\varepsilon \in \mathbb{Q}$, $\varepsilon > 0$ existiert ein $n_0 \in \mathbb{N}$ mit:

$$|x_n| < \varepsilon \quad \text{für alle } n \geq n_0.$$

- (ii) Zwei Fundamentalfolgen $x = \langle x_n \mid n \in \mathbb{N} \rangle$ und $y = \langle y_n \mid n \in \mathbb{N} \rangle$ haben den *gleichen Limes*, falls $\langle x_n - y_n \mid n \in \mathbb{N} \rangle$ eine Nullfolge ist.

Man kann nun die reellen Zahlen einfach als die Menge aller Fundamentalfolgen in $\mathbb{Q}$ definieren. Zwei reelle Zahlen gelten dann als *gleich*, falls sie den gleichen Limes haben. (Dies ist eine Gleichheit im Sinne einer Äquivalenzrelation.) Die rationalen Zahlen lassen sich in die reellen Zahlen einbetten, indem $q \in \mathbb{Q}$ mit der Fundamentalfolge identifiziert wird, die konstant den Wert q annimmt.

Diese Konstruktion der reellen Zahlen stammt von Georg Cantor aus dem gleichen Jahr wie die Konstruktion von Dedekind (1872). Unabhängig von Cantor wurde sie bereits früher von Charles Méray (1835 – 1911) durchgeführt (1869, 1872).

Heute wird üblicherweise in dieser Konstruktion eine reelle Zahl nicht als Fundamentalfolge, sondern als Äquivalenzklasse $x/\sim$ von Fundamentalfolgen definiert, wobei $x \sim y$, falls x und y den gleichen Limes haben. Dies hat lediglich die Konsequenz, daß der Gleichheitsbegriff auf den reellen Zahlen zur gewohnten Identität wird.

Für diese und weitere Konstruktionsmöglichkeiten und Charakterisierungen der reellen Zahlen verweisen wir den Leser auf [Deiser 2007] und [Rautenberg 2007]. Beide Texte berücksichtigen auch geschichtliche Entwicklungen.

Die Suslin-Hypothese

Bereits eine kleinere Variante des Trios „unbeschränkt, vollständig, separabel“ wirft große Probleme auf. Hierzu vorab zwei natürliche Begriffe.

Definition (*Intervalle in linearen Ordnungen*)

Sei $\langle M, < \rangle$ eine lineare Ordnung.

Ein $I \subseteq M$ heißt ein *Intervall* in $\langle M, < \rangle$, falls für alle $a, b \in I$ gilt:

Ist $c \in M$ und $a < c < b$, so ist $c \in I$.

Wie für reelle Intervalle sind „offen“, „geschlossen“ und $]a, b[$, $[a, b[$, $]a, b]$, $[a, b]$ definiert.

Definition (*Antiketten-Bedingung*)

Sei $\langle M, < \rangle$ eine lineare Ordnung.

$\langle M, < \rangle$ erfüllt die (*abzählbare*) *Antiketten-Bedingung*, falls jede Menge von paarweise disjunkten offenen Intervallen von M abzählbar ist.

paarweise disjunkte offene Intervalle in M

Übung

$\langle \mathbb{R}, < \rangle$ erfüllt die Antiketten-Bedingung.

Allgemeiner: Ist $\langle M, < \rangle$ separabel, so erfüllt $\langle M, < \rangle$ die Antiketten-Bedingung.

[Jedes nichtleere offene Intervall in $\mathbb{R}$ enthält eine rationale Zahl.]

Die Antiketten-Bedingung ist also eine Abschwächung der Separabilität. Die Frage ist, ob die Antiketten-Bedingung die Separabilität in der Charakterisierung des Typs θ ersetzen kann. Dies führt zur Suslin-Hypothese [Suslin, 1920].

Suslin-Hypothese

Sei $\langle M, < \rangle$ eine lineare Ordnung. Es gelte:

- (i) $\langle M, < \rangle$ ist unbeschränkt,
- (ii) $\langle M, < \rangle$ ist vollständig,
- (iii) $\langle M, < \rangle$ ist dicht und erfüllt die Antiketten-Bedingung.

Dann gilt $\theta = \text{o. t.}(\langle M, < \rangle)$, d. h. $\langle M, < \rangle$ und $\langle \mathbb{R}, < \rangle$ sind ähnlich.

Man kann nun versuchen, den Beweis der Charakterisierung von θ zu verfeinern. Oder man kann versuchen, ein Gegenbeispiel zu konstruieren, d. h. man sucht nach einer linearen Ordnung, die (i) – (iii) erfüllt, die aber keine abzählbare dichte Teilmenge enthält ...

Beide Ansätze sind hoffnungslos! Denn es zeigt sich: Die Suslin-Hypothese ist – genau wie die Kontinuumshypothese – weder beweisbar noch widerlegbar (im Rahmen der üblichen Mathematik). Der Beweis verläuft wie der Beweis der Unabhängigkeit von (CH) in zwei Teilen. Im Gödelschen Modell ist die Suslin-Hypothese falsch (siehe [Jensen 1968]). Mit der forcing-Methode von Cohen kann man sowohl Modelle konstruieren, in denen die Suslinhypothese richtig ist (siehe [Solovay/Tennenbaum 1971]), als auch solche, in denen sie falsch ist (siehe [Jech 1967, Tennenbaum 1968]).

Der Kern der Suslin-Hypothese ist der Zusammenhang von „separabel“ und „Antiketten-Bedingung“ für dichte lineare Ordnungen. Die Vollständigkeit und Unbeschränktheit der Ordnung ist nicht wesentlich:

Übung

Die folgenden Aussagen sind äquivalent:

- (i) Die Suslin-Hypothese ist richtig.
- (ii) Jede dichte lineare Ordnung, die die Antiketten-Bedingung erfüllt, ist separabel.

[zu (i) $\cap$ (ii): Wir zeigen $\neg(ii) \cap \neg(i)$. Sei $\langle M, < \rangle$ ein Gegenbeispiel zu (ii). Entfernung von Endpunkten und Dedekind-Vervollständigung von $\langle M, < \rangle$ liefert ein Gegenbeispiel zu (i); die Vervollständigung ist nicht separabel.]

Die einfache Aussage (ii) über lineare Ordnungen ist also ebenfalls im Rahmen der üblichen Mathematik nicht entscheidbar.

Wir beenden dieses Kapitel mit einem Kommentar Cantors zur philosophischen Diskussion um den Begriff eines Kontinuums und zum Verständnis der reellen Zahlen in der zweiten Hälfte des 19. Jahrhunderts. Damals wurde ein strenger Aufbau der Analysis zur dringenden Notwendigkeit, und die mangelnde Fundierung der reellen Zahlen war eine schwerwiegende Lücke dieses Unternehmens. Der folgende Auszug ist aber alleine schon wegen der darin dargelegten verblüffend modernen Sicht der Zeit von Interesse.

Georg Cantor über die mathematische Behandlung der reellen Zahlen

„Der Begriff des ‚Kontinuums‘ hat in der Entwicklung der Wissenschaften überall nicht nur eine bedeutende Rolle gespielt, sondern auch stets die größten Meinungsverschiedenheiten und sogar heftige Streitigkeiten hervorgerufen. Dies liegt vielleicht daran, daß die ihm zu Grunde liegende Idee in ihrer Erscheinung bei den Dissentierenden einen verschiedenen Inhalt aus dem Grunde angenommen hat, weil ihnen die genaue und vollständige Definition des Begriffs nicht überliefert worden war; vielleicht aber auch, und dies ist mir das Wahrscheinlichste, ist die Idee des Kontinuums schon von denjenigen Griechen, welche sie zuerst gefaßt haben mögen, nicht mit der Klarheit und Vollständigkeit gedacht worden, welche erforderlich gewesen wäre, um die Möglichkeit verschiedener Auffassungen seitens der Nachfolger auszuschließen. So sehen wir, daß Leukipp,

Demokrit und Aristoteles das Kontinuum als ein Kompositum betrachteten, welches ex partibus sine fine divisibilibus besteht, dagegen Epikur und Lukretius dasselbe aus ihren Atomen, als endlichen Dingen zusammensetzen, woraus nachmals ein großer Streit unter den Philosophen entstanden ist, von denen einige dem Aristoteles, andere dem Epikur gefolgt sind; andere wieder statuierten, um diesem Streit fern zu bleiben, mit Thomas von Aquino, daß das Kontinuum weder aus unendlich vielen, noch aus einer endlichen Anzahl von Teilen, sondern aus *gar keinen* Teilen bestehe; diese letztere Meinung scheint mir weniger eine Sacherklärung als das stillschweigende Bekenntnis zu enthalten, daß man der Sache nicht auf den Grund gekommen ist und es vorzieht, ihr vornehm aus dem Wege zu gehen. Hier sehen wir den *mittelalterlich-scholastischen Ursprung* einer Ansicht, die wir noch heutigen Tages vertreten finden, wonach das Kontinuum ein unzerlegbarer Begriff oder auch, wie andere sich ausdrücken, eine reine aprioristische Anschauung sei, die kaum einer Bestimmung durch Begriffe zugänglich wäre; jeder arithmetische *Determinationsversuch* dieses *Mysteriums* wird als ein unerlaubter Eingriff angesehen und mit gehörigem Nachdruck zurückgewiesen; schüchterne Naturen empfangen dabei den Eindruck, als ob es sich bei dem ‚Kontinuum‘ nicht um einen *mathematisch-logischen Begriff* sondern viel eher um ein *religiöses Dogma* handle.

Mir liegt es sehr fern, diese Streitfragen wieder heraufzubeschwören, auch würde mir zu einer genaueren Besprechung derselben in diesem engen Rahmen der Raum fehlen; ich sehe mich nur verpflichtet, den Begriff des Kontinuums, so logisch-nüchtern wie ich ihn auffassen muß und in der Mannigfaltigkeitslehre ihn brauche, hier möglichst kurz und auch nur mit Rücksicht auf die *mathematische* Mengenlehre zu entwickeln. Diese Bearbeitung ist mir aus dem Grunde nicht leicht geworden, weil unter den Mathematikern, auf deren Autorität ich mich gern berufe, kein Einziger sich mit dem Kontinuum in dem Sinne genauer beschäftigt hat, wie ich es hier nötig habe ...

Zunächst habe ich zu erklären, daß meiner Meinung nach die Heranziehung des *Zeitbegriffs* oder der *Zeitanschauung* bei Erörterung des viel ursprünglicheren und allgemeineren Begriffs des Kontinuums nicht in der Ordnung ist; die *Zeit* ist meines Erachtens eine Vorstellung, die zu ihrer deutlichen Erklärung den von ihr unabhängigen Kontinuitätsbegriff zur Voraussetzung hat und sogar mit Zuhilfenahme desselben weder objektiv als eine Substanz, noch subjektiv als eine notwendige apriorische Anschauungsform aufgefaßt werden kann, sondern nichts anderes als ein *Hilfs- und Beziehungsbegriff*, durch welchen die Relation zwischen verschiedenen in der Natur vorkommenden und von uns wahrgenommenen Bewegungen festgestellt wird. So etwas wie *objektive* oder *absolute Zeit* kommt in der Natur nirgends vor und es kann daher auch nicht die *Zeit* als Maß der Bewegung, viel eher könnte diese als Maß der *Zeit* angesehen werden, wenn nicht dem Letzteren entgegenstände, daß die *Zeit* selbst in der bescheidenen Rolle einer *subjektiv notwendigen apriorischen* Anschauungsform es zu keinem ersprißlichen, unangefochtenen Gedeihen hat bringen können, obgleich ihr seit Kant die *Zeit* dazu nicht gefehlt haben würde.“

(Georg Cantor 1883b, „Über unendliche lineare Punktmannigfaltigkeiten. V“)

11. Der Satz von Cantor-Bendixson

Wir beantworten in diesem und im nächsten Kapitel die beiden offen gebliebenen Fragen aus dem zweiten Kapitel:

- (1) Wieviele Ableitungen $P^{(\alpha)}$ braucht man, um von einem $P \subseteq \mathbb{R}$ zu einer perfekten Menge zu gelangen, d. h. zu einem $P^{(\alpha)}$ mit $P^{(\alpha)} = P^{(\alpha+1)}$?
- (2) Welche Mächtigkeiten sind für die abgeschlossenen Mengen möglich?

Für eine beliebige Teilmenge P von $\mathbb{R}$ hatten wir rekursiv die Ableitungen $P^{(\alpha)}$ für Ordinalzahlen α definiert. Zur Erinnerung:

Definition (die Ableitungen $P^{(\alpha)}$ einer Punktmenge)

Sei $P \subseteq \mathbb{R}$. Dann ist die α -te Ableitung $P^{(\alpha)}$ von P für Ordinalzahlen α rekursiv definiert durch:

$$P^{(0)} = P,$$

$$P^{(\alpha+1)} = P^{(\alpha)'} \quad \text{für } \alpha \text{ Ordinalzahl,}$$

$$P^{(\lambda)} = \bigcap_{\alpha < \lambda} P^{(\alpha)} \quad \text{für } \lambda \text{ Limesordinalzahl.}$$

Nach den Ergebnissen des zweiten Kapitels wissen wir, daß die Ableitungen ab dem ersten Schritt eine $\subseteq$ -absteigende Folge aus abgeschlossenen Mengen bildet:

Satz

Sei $P \subseteq \mathbb{R}$. Dann gilt für alle Ordinalzahlen $\alpha \neq 0$:

- (i) $P^{(\alpha)}$ ist abgeschlossen,
- (ii) $P^{(\alpha)} \supseteq P^{(\alpha+1)}$.

Wir haben also $P^{(1)} \supseteq P^{(2)} \supseteq \dots \supseteq P^{(\alpha)} \supseteq P^{(\alpha+1)} \supseteq \dots$

Gilt $P^{(\alpha)} = P^{(\alpha+1)}$ für eine Ordinalzahl α , so ist der Ableitungsprozeß ab α konstant gleich $P^{(\alpha)}$. Mit $P^{(\alpha)}$ ist dann eine perfekte Menge gefunden, die wir als den *perfekten Kern* aller abgeschlossenen Mengen der Ableitungsfolge bezeichnen:

Definition (perfekter Kern)

Sei $P \subseteq \mathbb{R}$ abgeschlossen, und sei α eine Ordinalzahl mit $P^{(\alpha)} = P^{(\alpha+1)}$.

Dann heißt $P^{(\alpha)}$ der *perfekte Kern* von P .

Wir wollen nun zeigen, daß wir den perfekten Kern einer abgeschlossenen Menge in abzählbar vielen Schritten erreichen. Hierzu beweisen wir einen allgemeinen Satz über die Länge von $\subsetneq$ -absteigenden Folgen abgeschlossener Teilmengen von $\mathbb{R}$.

Absteigende Folgen abgeschlossener Mengen

Definition (*$\subset$ -absteigende Folgen von Teilmengen von $\mathbb{R}$*)

Eine Folge $\langle X_\alpha \mid \alpha < \beta \rangle$ von Mengen heit *$\subset$ -absteigend* (oder *strikt absteigend bzgl. der Inklusion*), falls gilt:

Fr alle $\alpha, \alpha' < \beta$ mit $\alpha < \alpha'$ gilt $X_{\alpha'} \subset X_\alpha$.

Analog sind *$\subseteq$ -absteigend*, *$\subset$ -aufsteigend* ($X_\alpha \subset X_{\alpha'}$ fr $\alpha < \alpha' < \beta$) und *$\subseteq$ -aufsteigend* definiert.

Allgemein knnen $\subset$ -absteigende Folgen von linearen Punktmengen berabzhlbare Lnge haben:

bung

Es gibt eine $\subset$ -absteigende Folge $\langle X_\alpha \mid \alpha < \omega_1 \rangle$ von Teilmengen von $\mathbb{R}$.

[Sei $f: W(\omega_1) \rightarrow \mathbb{R}$ injektiv. Wir setzen $X_\alpha = \mathbb{R} - f''W(\alpha)$ fr $\alpha < \omega_1$.]

Fordern wir aber zustzlich, da die Folge aus *abgeschlossenen* Teilmengen von $\mathbb{R}$ gebildet ist, so ist ein derartig langsamer Abstieg nicht mehr mglich, wie wir nun zeigen wollen. Entscheidend verwenden wir hierbei die folgende Konsequenz der Separabilitt von $\mathbb{R}$: Jedes offene reelle Intervall enthlt ein offenes Intervall mit rationalen Endpunkten.

Definition (*rationales Intervall*)

Sei $I =]a, b[$ ein offenes reelles Intervall mit $a, b \in \mathbb{Q}$.

Dann heit I ein *offenes rationales Intervall*.

bung

- (i) Sei $I =]a, b[$ ein nichtleeres reelles Intervall, und sei $x \in I$.

Dann existiert ein nichtleeres rationales Intervall

$J =]a', b'[$ mit $x \in J$ und $J \subseteq I$.

- (ii) Ist $Q \subseteq \mathbb{R}$ offen, so ist $Q = \bigcup \{ I \mid I \subseteq Q, I \text{ rationales Intervall} \}$.

- (iii) Die Menge der rationalen offenen Intervalle ist abzhlbar.

[$\mathbb{Q} \times \mathbb{Q}$ ist abzhlbar.]

Ist $\langle P_\alpha \mid \alpha < \beta \rangle$ eine $\subset$ -absteigende Folge von Teilmengen von $\mathbb{R}$, und enthlt $P_\alpha - P_{\alpha+1}$ fr alle $\alpha < \beta$ eine rationale Zahl, so ist β offenbar abzhlbar. Die Differenz $A - B$ zweier abgeschlossener Teilmengen $B \subset A$ enthlt aber i. a. keine rationale Zahl. Jedoch fhrt eine Verfeinerung dieser Idee zum Ziel: Es gibt ein offenes rationales Intervall, das A trifft, aber zu B disjunkt ist. Diese Beobachtung bildet das Herz des folgenden Beweises.

Satz (über die Länge absteigender Folgen von abgeschlossenen Mengen)

Sei $\langle P_\alpha \mid \alpha < \beta \rangle$ eine $\subset$ -absteigende Folge abgeschlossener Teilmengen von $\mathbb{R}$.

Dann ist β abzählbar.

Beweis

Sei $I_0, I_1, I_2, \dots$ eine Aufzählung aller offenen nichtleeren rationalen Intervalle in $\mathbb{R}$. Für $\alpha < \beta$ definieren wir $N_\alpha \subseteq \mathbb{N}$ durch

$$N_\alpha = \{n \in \mathbb{N} \mid I_n \cap P_\alpha \neq \emptyset\}.$$

Dann gilt:

(+) $\langle N_\alpha \mid \alpha < \beta \rangle$ ist $\subset$ -absteigend.

Beweis von (+)

Offenbar gilt $N_{\alpha'} \subseteq N_\alpha$ für alle $\alpha < \alpha' < \beta$.

Es genügt zu zeigen: $N_{\alpha+1} \subseteq N_\alpha$ für alle α mit $\alpha + 1 < \beta$.

Sei also $\alpha + 1 < \beta$. Wegen $P_{\alpha+1} \subset P_\alpha$ existiert ein $x \in P_\alpha - P_{\alpha+1}$.

Dann gilt $N_x = \{n \in \mathbb{N} \mid x \in I_n\} \subseteq N_\alpha$.

Wäre nun $P_{\alpha+1} \cap I_n \neq \emptyset$ für alle $n \in N_x$, so wäre x ein

Häufungspunkt von $P_{\alpha+1}$. Aber $x \notin P_{\alpha+1}$, im Widerspruch zur Abgeschlossenheit von $P_{\alpha+1}$.

Also existiert ein $n \in N_x \subseteq N_\alpha$ mit $n \notin N_{\alpha+1}$. Dies zeigt (+).

Also ist $\langle N_\alpha \mid \alpha < \beta \rangle$ eine $\subset$ -absteigende Folge von Teilmengen von $\mathbb{N}$.

Dann ist aber β abzählbar, denn die Funktion $g: \text{dom}(g) \rightarrow \mathbb{N}$ mit

$$g(\alpha) = \min(N_\alpha - N_{\alpha+1}) \quad \text{für alle } \alpha \text{ mit } \alpha + 1 < \beta$$

– ist injektiv.

Übung

Zeigen Sie, daß auch $\subset$ -aufsteigende Folgen abgeschlossener Mengen abzählbare Länge haben. Ebenso sind $\subset$ -aufsteigende oder $\subset$ -absteigende Folgen offener Mengen abzählbar.

Hausdorff (1914): „Eine auf- oder absteigende wohlgeordnete Menge verschiedener Gebiete [offener Mengen] oder abgeschlossener Mengen ist höchstens abzählbar.“

Aus dem Satz ergibt sich, daß die Folge der Ableitungen einer abgeschlossenen Menge nach abzählbar vielen Schritten in einer perfekten Menge terminiert, denn andernfalls hätten wir mit $P^{(1)} \supset \dots \supset P^{(\alpha)} \supset P^{(\alpha+1)} \supset \dots$ eine überabzählbare $\subset$ -absteigende Folge abgeschlossener Mengen vor uns.

Damit ist die erste Frage beantwortet: Abzählbar viele Schritte genügen, um in der Folge der Ableitungen einen stabilen Zustand zu erreichen.

Der Satz von Cantor-Bendixson macht nun zudem eine Aussage über die während der Ableitung verlorengegangenen Punkte.

Die Cantor-Bendixson-Zerlegung

Satz (Mächtigkeit der isolierten Punkte einer Menge)

Sei $P \subseteq \mathbb{R}$. Dann ist die Menge $R = P - P'$ der isolierten Punkte von P abzählbar.

Beweis

Sei wieder $I_0, I_1, I_2, \dots$ eine Aufzählung aller offenen nichtleeren rationalen Intervalle in $\mathbb{R}$.

Für alle $x \in R$ sei

$f(x) = \text{„das kleinste } n \in \mathbb{N} \text{ mit } x \in I_n \text{ und } I_n \cap P = \{x\}\text{“}.$

$f(x)$ ist wohldefiniert, da x ein isolierter Punkt von P ist.

Dann ist $f : R \rightarrow \mathbb{N}$ injektiv, denn für $x, y \in R$ mit $f(x) = f(y)$ ist

$$\{x\} = I_{f(x)} \cap P = I_{f(y)} \cap P = \{y\},$$

also $x = y$.

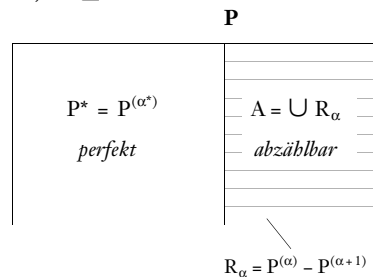
– Also ist R abzählbar.

Damit erhalten wir nun den zusammenfassenden Satz:

Satz (Satz von Cantor-Bendixson)

Sei $P \subseteq \mathbb{R}$ abgeschlossen. Dann existieren $A, P^* \subseteq P$ mit:

- (i) $P^* \cap A = \emptyset, P^* \cup A = P,$
- (ii) $P^* = P^{(\alpha)}$ für ein abzählbares $\alpha,$
- (iii) P^* ist perfekt, A ist abzählbar.



Beweis

Wir setzen:

$\alpha^* = \text{„das kleinste } \alpha \text{ mit } P^{(\alpha)} = P^{(\alpha+1)}\text{“},$

$P^* = P^{(\alpha^*)}.$

Nach dem Satz über die Länge $\subset$ -absteigender Folgen abgeschlossener Mengen existiert α^* und ist abzählbar. Offenbar ist $P^* \subseteq P$ perfekt.

Für $\alpha < \alpha^*$ sei $R_{\alpha} = P^{(\alpha)} - P^{(\alpha+1)}$ die Menge der isolierten Punkte von $P^{(\alpha)}$.

Sei dann

$$A = \bigcup_{\alpha < \alpha^*} R_{\alpha}.$$

Dann ist A als abzählbare Vereinigung abzählbarer Mengen abzählbar.

– Offenbar gilt $P = P^* \cup A$ und $P^* \cap A = \emptyset$.

Definition (*Cantor-Bendixson-Zerlegung*)

Sei $P \subseteq \mathbb{R}$ abgeschlossen.

Dann heißt die im obigen Beweis konstruierte Zerlegung

$$P = P^* \cup A$$

die *Cantor-Bendixson-Zerlegung* von P .

A heißt die Menge der *verborgen isolierten Punkte* von P .

Cantor (1884b): „Theorem E'. Ist P eine abgeschlossene Punktmenge von höherer als der ersten Mächtigkeit [P überabzählbar], so zerfällt P und zwar nur auf eine Weise in eine [nichtleere] perfekte Menge S und eine Menge von der ersten Mächtigkeit [hier: höchstens abzählbar] R , so daß

$$P \equiv R + S$$

und es existiert eine kleinste der ersten oder zweiten Zahlenklasse zugehörige [abzählbare] Zahl α , so daß $P^{(\alpha)}$ gleich S wird.“

Die Eindeutigkeit der Zerlegung einer abgeschlossenen Menge in einen perfekten und einen abzählbaren Teil zeigen wir unten. Hier halten wir noch ein einfaches Korollar fest:

Korollar

Sei $P \subseteq \mathbb{R}$, und sei $P^{(\alpha)}$ abzählbar für ein α .

Dann ist P abzählbar.

Beweis

Nichttrivial ist nur der Fall $\alpha \geq 1$.

Sei $P^{(1)} = P^* \cup A$ die Cantor-Bendixson-Zerlegung der abgeschlossenen Menge $P^{(1)}$. Dann ist P^* abzählbar, denn P^* ist Teilmenge jeder Ableitung $P^{(\alpha)}$ von P .

Wegen A abzählbar ist also $P^{(1)}$ abzählbar.

– Aber $R = P - P^{(1)}$ ist abzählbar, also ist $P = P^{(1)} \cup R$ abzählbar.

Cantor (1884b): „Theorem B. Ist α irgendeine Zahl der ersten oder zweiten Zahlenklasse [α abzählbar] und $P \dots$ eine Punktmenge von solcher Beschaffenheit, daß:

$$P^{(\alpha)} \equiv 0,$$

so ist $P^{(1)}$ sowohl wie auch P von der ersten Mächtigkeit [abzählbar unendlich], es sei denn, daß P resp. $P^{(1)}$ endliche Mengen sind.“

Aus dem Satz von Cantor-Bendixson folgt, daß jede überabzählbare abgeschlossene Menge immer einen überabzählbaren perfekten Kern enthält. Umgekehrt stellt sich die Frage:

Kann eine abzählbare abgeschlossene Menge einen nichtleeren Kern haben?

Dieser Frage wollen wir uns nun zuwenden.

Reduktible Teilmengen von $\mathbb{R}$

Definition (*reduktibel*)

Sei $P \subseteq \mathbb{R}$. P heißt *reduktibel*,
falls eine Ordinalzahl γ existiert mit $P^{(\gamma)} = \emptyset$.

 P
 P'
 $\dots$
 $P^{(\alpha)}$
 $\dots$
 $\dots$
 $P^{(\gamma)} = \emptyset$

Obige Frage lautet also: Ist jede abgeschlossene abzählbare Menge reduktibel? Oder: Kann eine nichtleere perfekte Menge abzählbar sein? Jede nichtleere perfekte Menge $P \subseteq \mathbb{R}$ ist offenbar unendlich, da jedes $x \in P$ ein Häufungspunkt von P ist. Mit einem Argument, das dem ersten Cantorschen Beweis der Überabzählbarkeit der reellen Zahlen verwandt ist, zeigen wir nun, daß nichttriviale perfekte Mengen sogar überabzählbar sind.

Satz (*Überabzählbarkeit der perfekten Mengen*)

Sei $P \subseteq \mathbb{R}$ perfekt und nichtleer. Dann ist P überabzählbar.

Beweis

Sei $f: \mathbb{N} \rightarrow P$ injektiv, und sei $x_n = f(n)$ für $n \in \mathbb{N}$.

Wir konstruieren ein $x^* \in P$, das von allen x_n , $n \in \mathbb{N}$, verschieden ist.

Dann kann f keine Bijektion sein, und es folgt, daß P überabzählbar ist.

Für alle $n \in \mathbb{N}$ können wir annehmen:

(+) Für alle $\varepsilon > 0$, $\varepsilon \in \mathbb{R}$, ist $(U_\varepsilon(x_n) - \{x_n\}) \cap \text{rng}(f) \neq \emptyset$.

Andernfalls ist f sicher nicht bijektiv: Denn $(U_\varepsilon(x_n) - \{x_n\}) \cap P$ ist unendlich, da x_n ein Häufungspunkt von P ist.

Wir konstruieren nun rekursiv:

- eine konvergente Teilfolge $x_{i_0}, x_{i_1}, x_{i_2} \dots$ von $x_0, x_1, x_2, \dots$,
wobei $i_0 < i_1 < i_2 < \dots$,
- eine Nullfolge $\varepsilon_0 > \varepsilon_1 > \varepsilon_2 > \dots$ von reellen Zahlen $\varepsilon_n > 0$.

Sei $i_0 = 0$, $\varepsilon_0 = 1$.

Seien i_n und ε_n konstruiert für ein $n \in \mathbb{N}$. Wir setzen:

$i_{n+1} = \text{„das kleinste } m \in \mathbb{N} \text{ mit } x_m \in U_{\varepsilon_n}(x_{i_n}) - \{x_{i_n}\} \text{“},$

$\varepsilon_{n+1} = 1/2 \cdot \min(\sigma_n, \varepsilon_n - \sigma_n)$, wobei $\sigma_n = |x_{i_{n+1}} - x_{i_n}|$.

(Die Idee ist, daß $U_{\varepsilon_{n+1}}(x_{i_{n+1}})$ samt Rand ganz in $U_{\varepsilon_n}(x_{i_n}) - \{x_{i_n}\}$ liegt.)

Nach (+) ist i_n wohldefiniert für alle $n \in \mathbb{N}$, und es gilt $i_0 < i_1 < i_2 < \dots$ (!).

Weiter ist $\varepsilon_{n+1} < 1/2 \varepsilon_n$ für alle $n \in \mathbb{N}$, also $\lim_{n \rightarrow \infty} \varepsilon_n = 0$.

Nach Konstruktion ist $\langle x_{i_n} \mid n \in \mathbb{N} \rangle$ eine Cauchyfolge, also existiert

$x^* = \lim_{n \rightarrow \infty} x_{i_n}$.

Dann ist $x^* \in P$, da x^* ein Häufungspunkt von P ist.

- Aber nach Konstruktion gilt $x^* \neq x_n$ für alle $n \in \mathbb{N}$.

Cantor (1884b): „Theorem A. Eine ... Punktmenge P kann, wenn sie von der ersten Mächtigkeit [abzählbar unendlich] ist, nie eine perfekte Punktmenge sein.“

Mit diesem Ergebnis können wir die Frage nach der Reduktibilität abzählbarer abgeschlossener Mengen positiv beantworten:

Korollar (*abgeschlossene abzählbare Mengen sind reduktibel*)

Sei $P \subseteq \mathbb{R}$ abgeschlossen und abzählbar.

Dann existiert eine abzählbare Ordinalzahl α mit $P^{(\alpha)} = \emptyset$.

Beweis

Sei α eine abzählbare Ordinalzahl, für die $P^{(\alpha)}$ perfekt ist.

Wegen P abgeschlossen ist $P^{(\alpha)} \subseteq P$, also ist $P^{(\alpha)}$ perfekt und abzählbar.

– Aber nichtleere perfekte Mengen sind überabzählbar, also ist $P^{(\alpha)} = \emptyset$.

Cantor (1884b): „Theorem C'. Ist P irgend eine abgeschlossene Punktmenge von der ersten Mächtigkeit [abzählbar unendlich], so gibt es immer eine kleinste der ersten oder zweiten Zahlenklasse zugehörige Zahl α [ein abzählbares α], so daß $P^{(\alpha)}$ gleich Null ist, oder was dasselbe heißen soll, solche Mengen sind immer reduktibel.“

Übung

Sei $P \neq \emptyset$ abgeschlossen und abzählbar, und sei α minimal mit $P^{(\alpha)} = \emptyset$.

Dann ist α eine Nachfolgerordinalzahl.

Explizit halten wir die folgende Umformulierung fest:

Korollar

Sei $P \subseteq \mathbb{R}$ nichtleer, abgeschlossen und abzählbar.

Dann hat P einen isolierten Punkt.

Eine weitere Folge der Überabzählbarkeit nichtleerer perfekter Mengen ist die Eindeutigkeit der Cantor-Bendixson-Zerlegung.

Korollar (*Eindeutigkeit der Cantor-Bendixson-Zerlegung*)

Sei $P \subseteq \mathbb{R}$, und sei $P = P_1 \cup A_1 = P_2 \cup A_2$ mit

(i) P_1, P_2 perfekt,

(ii) A_1, A_2 abzählbar,

(iii) $P_1 \cap A_1 = P_2 \cap A_2 = \emptyset$.

Dann gilt $P_1 = P_2$ und $A_1 = A_2$.

Beweis

Annahme $P_1 \neq P_2$. Sei dann o. E. $x \in P_1, x \notin P_2$. Wegen P_2 abgeschlossen existiert ein Intervall $I =]a, b[$ mit $x \in I$ und $I \subseteq \mathbb{R} - P_2$.

Weiter existiert ein Intervall $I' = [a', b']$ mit $a < a' < x < b' < b$.

Dann ist $P_1 \cap I'$ nach evtl. Entfernung von a' und b' perfekt (!) und nichtleer. Also ist $P_1 \cap I'$ überabzählbar. Aber es gilt $P_1 \cap I' \subseteq P \cap I \subseteq A_2$, und damit ist – A_2 überabzählbar. *Widerspruch!* Also gilt $P_1 = P_2$ und weiter $A_1 = A_2$.

Wir fassen die wesentlichen Resultate noch einmal in einem Satz zusammen.

Satz (*Fundamentalsatz über Ableitungen*)

Sei $P \subseteq \mathbb{R}$. Dann gilt:

- (A) Ist $P^{(\alpha)}$ abzählbar für eine Ordinalzahl α , so ist P abzählbar.
- (B) Ist P abzählbar und abgeschlossen, so ist $P^{(\alpha)} = \emptyset$ für ein abzählbares α .
- (C) Ist P überabzählbar und abgeschlossen, so existiert ein perfektes $P^* \subseteq P$, für das $P - P^*$ abzählbar ist. Ein solches P^* ist eindeutig bestimmt, und es gilt $P^* = P^{(\alpha)}$ für ein abzählbares α .

Die im Hinblick auf (C) an dieser Stelle natürliche Frage, ob jede überabzählbare lineare Punktmenge eine nichtleere perfekte Teilmenge besitzt, behandeln wir im nächsten Kapitel.

Kondensationspunkte

Wir wollen nun noch eine andere Konstruktion der Cantor-Bendixson-Zerlegung geben, die den perfekten Kern einer abgeschlossenen Punktmenge unmittelbar aus P extrahiert, dafür aber die durch den Ableitungsprozeß aufgedeckte Struktur einer Punktmenge nicht ans Licht bringt.

Hierzu betrachten wir zunächst die folgende Verschärfung des Begriffs eines Häufungspunktes einer Punktmenge (Lindelöf, 1905):

Definition (*Kondensationspunkte einer Punktmenge*)

Sei $P \subseteq \mathbb{R}$. Ein $x \in \mathbb{R}$ heißt ein *Kondensationspunkt* von P , falls gilt:

Für alle $\varepsilon > 0$, $\varepsilon \in \mathbb{R}$, ist $U_\varepsilon(x) \cap P$ überabzählbar.

Wir setzen:

$\text{cp}(P) = \{x \in \mathbb{R} \mid x \text{ ist Kondensationspunkt von } P\}$.

[cp steht hier für engl. *condensation points*.]

Zum Vergleich: x ist Häufungspunkt von P , falls gilt:
Für alle $\varepsilon > 0$ ist $U_\varepsilon(x) \cap P$ unendlich.

Satz

Sei $P \subseteq \mathbb{R}$. Dann gilt:

- (i) $P - \text{cp}(P)$ ist abzählbar.
- (ii) $\text{cp}(P)$ ist perfekt.

Beweis

zu (i): Sei $A = P - \text{cp}(P)$. Dann gilt für alle $x \in P$:

(+) $x \in A$ gdw es existiert ein offenes rationales Intervall I mit:
 $x \in I$ und $I \cap P$ ist abzählbar.

Also gilt:

$$A = \bigcup \{ I \cap P \mid I \text{ ist ein offenes rationales Intervall und } I \cap P \text{ ist abzählbar} \}.$$

Also ist A abzählbare Vereinigung abzählbarer Mengen, und damit abzählbar.

zu (ii): $\text{cp}(P)$ ist abgeschlossen, denn ein Häufungspunkt von Kondensationspunkten von P ist wieder ein Kondensationspunkt (vgl. den Beweis von „ P' abgeschlossen“):

Ist x ein Häufungspunkt von $\text{cp}(P)$ und ist $\varepsilon > 0$, so existiert ein $y \in \text{cp}(P)$ und ein $\delta > 0$ mit $U_\delta(y) \subseteq U_\varepsilon(x)$.

Dann ist $U_\delta(y) \cap P$ überabzählbar und eine Teilmenge von $U_\varepsilon(x) \cap P$.

Also ist $U_\varepsilon(x) \cap P$ überabzählbar für alle $\varepsilon > 0$, also $x \in \text{cp}(P)$.

Wir zeigen noch, daß $\text{cp}(P)$ keine isolierten Punkte besitzt.

Sei hierzu $x \in \text{cp}(P)$, und sei $\varepsilon > 0$.

Nach Definition von $\text{cp}(P)$ gilt dann:

(#) $U_\varepsilon(x) \cap P$ ist überabzählbar.

Wir müssen zeigen:

(##) $U_\varepsilon(x) \cap \text{cp}(P)$ ist unendlich.

– Aber $P - \text{cp}(P)$ ist abzählbar nach (i). Also folgt (##) aus (#).

Korollar

Sei $P \subseteq \mathbb{R}$ abgeschlossen, und sei $A = P - \text{cp}(P)$.

Dann ist $P = \text{cp}(P) \cup A$ die Cantor-Bendixson-Zerlegung von P .

Insbesondere ist also $\text{cp}(P)$ der perfekte Kern P^* von P .

Beweis

Wegen P abgeschlossen ist $\text{cp}(P) \subseteq P$. Nach dem Satz oben ist $\text{cp}(P)$ perfekt und A abzählbar. Aus der Eindeutigkeit der

– Cantor-Bendixson-Zerlegung folgt damit die Behauptung.

Übung

Sei $P \subseteq \mathbb{R}$ perfekt. Dann ist $\text{cp}(P) = P$.

Eine weitere Möglichkeit, den perfekten Kern aus einer Menge in einem einzigen Schritt zu gewinnen, gibt die folgende Übung. Dieser Ansatz geht auf Hausdorff zurück (vgl. [Hausdorff 1914, S. 221 ff] und das Zitat am Ende des Kapitels).

Übung

Sei $P \subseteq \mathbb{R}$ abgeschlossen, und sei

$$P^\Delta = \bigcup \{ A \subseteq P \mid A \text{ ist in sich dicht (d.h. } A \subseteq A') \}.$$

Dann ist P^Δ der perfekte Kern von P .

Der perfekte Kern von P ist also die größte in sich dichte Teilmenge von P .

Kohärenzen und relative Abgeschlossenheit

Die Folge der Ableitungen $P^{(\alpha)}$ mündet für beliebige $P \subseteq \mathbb{R}$ nach dem ersten Schritt in die Folge der Ableitungen einer abgeschlossenen Menge, nämlich in die Ableitungsfolge von P' . Ein feinerer, ebenfalls schon von Cantor untersuchter Ableitungsprozeß einer beliebigen Punktmenge ist der folgende:

Definition (*Kohärenzen und Adhärenzen*)

Sei $P \subseteq \mathbb{R}$. Dann ist die α -te Kohärenz P^α von P für Ordinalzahlen α rekursiv definiert durch:

$$P^0 = P,$$

$$P^{\alpha+1} = (P^\alpha)' \cap P^\alpha \quad \text{für } \alpha \text{ Ordinalzahl,}$$

$$P^\lambda = \bigcap_{\alpha < \lambda} P^\alpha \quad \text{für } \lambda \text{ Limesordinalzahl.}$$

$P^\alpha - P^{\alpha+1}$ heißt die α -te Adhärenz von P .

Die Begriffe der Kohärenz und Adhärenz hat Cantor 1885 eingeführt.

Die erste Adhärenz von P ist einfach die Menge der isolierten Punkte von P , und also abzählbar. Insgesamt werden bei der iterierten Kohärenzbildung also in jedem Schritt nur abzählbar viele Punkte abgespalten.

Für abgeschlossene Mengen P gilt $P^{(\alpha)} = P^\alpha$ für alle α . Dagegen gilt $\mathbb{Q}^{(1)} = \mathbb{R}$, während $\mathbb{Q}^1 = \mathbb{Q}$ gilt.

Im allgemeinen gilt $P^\alpha = P^{(\alpha)} \cap P$ nicht:

Übung

(i) Es gibt ein $P \subseteq \mathbb{R}$ mit $P^2 = \emptyset$, $P^{(2)} \subseteq P$, $P^{(2)} \neq \emptyset$.

(ii) Es gilt $P^\alpha \subseteq P^{(\alpha)} \cap P$ für alle Ordinalzahlen α und alle $P \subseteq \mathbb{R}$.

[zu (i): Sei $\langle x_\alpha \mid \alpha \leq \omega^2 \rangle$ strikt aufsteigend in $\mathbb{R}$.

Betrachte nun $P = \{ x_\alpha \mid \alpha < \omega^2, \alpha \text{ Nachfolger} \} \cup \{ x_{\omega^2} \}$.]

Ist also P reduktibel, so ist $P^\alpha = \emptyset$ für ein α . Die Umkehrung gilt nicht. Wir werden im nächsten Abschnitt ein $P \subseteq \mathbb{R}$ konstruieren, dessen erste Kohärenz $P^1 = P' \cap P$ leer ist, dessen erste Ableitung $P^{(1)} = P'$ aber nichtleer und perfekt ist.

Kohärenzen haben eine gewisse Abschlußeigenschaft. Hierzu ein auch anderswo nützlicher Begriff:

Definition (*abgeschlossen in P , $\mathcal{F}_P$*)

Sei $P \subseteq \mathbb{R}$. $A \subseteq P$ heißt (*relativ*) *abgeschlossen in P* , falls $A' \cap P \subseteq A$.

Wir setzen $\mathcal{F}_P = \{ A \subseteq P \mid A \text{ ist abgeschlossen in } P \}$.

Insbesondere ist jedes $P \subseteq \mathbb{R}$ abgeschlossen in sich selbst, und es gilt $\mathcal{F}_{\mathbb{R}} = \mathcal{F}$.

Die Idee ist: P ist blind gegenüber $\mathbb{R} - P$, insbesondere also gegenüber $P' - P$. Ein derart blindes P denkt, daß $A \subseteq P$ abgeschlossen ist, wenn jeder Häufungspunkt von A , den P sieht, bereits zu A gehört.

Übung

Für alle $A, P \subseteq \mathbb{R}$ sind äquivalent:

- (i) A ist abgeschlossen in P .
- (ii) Es gibt ein abgeschlossenes $B \subseteq \mathbb{R}$ mit $A = P \cap B$.

Übung

- (a) Ist $P \subseteq \mathbb{R}$ und $\mathcal{A} \subseteq \mathcal{F}_P$ nichtleer, so ist $\bigcap \mathcal{A} \in \mathcal{F}_P$.
- (b) Seien $A, B, C \subseteq \mathbb{R}$, und sei $A \subseteq B \subseteq C$. Dann gilt:
Ist A abgeschlossen in C , so ist A abgeschlossen in B .

Übung

Für alle $P \subseteq \mathbb{R}$ und alle Ordinalzahlen α gilt:
 P^α ist abgeschlossen in P .

Auch für die relative Abgeschlossenheit gilt der Satz über die Länge von absteigenden Folgen:

Satz (*über die Länge absteigender Folgen von relativ abgeschlossenen Mengen*)

Sei $P \subseteq \mathbb{R}$ und sei $\langle A_\alpha \mid \alpha < \beta \rangle$ eine $\subset$ -absteigende Folge von in P abgeschlossenen Teilmengen von P .
Dann ist β abzählbar.

Beweis

Sei $P \subseteq \mathbb{R}$, und seien $A, B \subseteq P$ abgeschlossen in P . Dann gilt:

Ist $A \subset B$, so ist $\text{cl}(A) \subset \text{cl}(B)$, d.h. $A \cup A' \subset B \cup B'$.

$\subseteq$ ist klar. Sei also $x \in B - A$. Dann ist $x \notin A'$, da A abgeschlossen in P ist. Also $x \notin A \cup A'$.

Folglich ist $\langle \text{cl}(A_\alpha) \mid \alpha < \beta \rangle$ eine $\subset$ -absteigende Folge von
— abgeschlossenen Teilmengen von $\mathbb{R}$. Also ist β abzählbar.

Auch der Prozeß der iterierten Kohärenzbildung erreicht also nach abzählbar vielen Schritten einen stabilen Zustand:

Satz (*abzählbare Terminierung der Kohärenzbildung*)

Sei $P \subseteq \mathbb{R}$. Dann existiert ein abzählbares α mit $P^\alpha = P^{\alpha+1}$.

Beweis

Andernfalls ist $\langle P^\alpha \mid \alpha < \omega_1 \rangle$ eine $\subset$ -absteigende Folge von in P

– abgeschlossenen Teilmengen von P , *Widerspruch*.

Definition (*letzte Kohärenz*)

Sei $P \subseteq \mathbb{R}$, und sei α derart, daß $P^\alpha = P^{\alpha+1}$ gilt.

Dann heißt P^α die *letzte Kohärenz von P* .

Die letzte Kohärenz einer Punktmenge läßt sich wieder wie folgt charakterisieren:

Übung

Sei $P \subseteq \mathbb{R}$, und sei $K \subseteq P$ die letzte Kohärenz von P .

Dann ist K die größte in sich dichte Teilmenge von P .

Residuen und reduzible Punktmenngen

Wir besprechen noch eine zur Kohärenzbildung einer Punktmenge verwandte Operation, die überaus originelle Hausdorffsche Residuenbildung.

Ist $P \subseteq \mathbb{R}$, so ist der Abschluß $\text{cl}(P) = P \cup P' = P \cup (P' - P)$ von P die kleinste abgeschlossene Menge, die P umfaßt. Im folgenden betrachten wir die beim Übergang von P zu $\text{cl}(P)$ neu hinzukommenden Punkte $\text{cl}(P) - P = P' - P$ genauer. Insbesondere betrachten wir wieder den Abschluß dieser Punkte, und die dabei zu $P' - P$ neu hinzukommenden Punkte, usw. Dabei entsteht ein interessanter Zickzackprozeß innerhalb von $\text{cl}(P)$, bei dem wir zwischen P und $\text{cl}(P) - P$ hin und her springen.

Für das folgende ist es nützlich, sich ein typisches $P \subseteq \mathbb{R}$ als eine fein zerstäubte Punktmenge vorzustellen (Puderzucker auf einem Kuchen). Dann ist $\text{cl}(P)$ sehr reich (die Oberfläche des Kuchens), und $\text{cl}(P) - P$ (der ungezuckerte Teil der Oberfläche) sieht ungefähr so aus wie P selbst.

Definition (*Hausdorff-Residuum, φ - und ψ -Operationen auf Punktmenngen*)

Sei $P \subseteq \mathbb{R}$. Wir definieren:

$$\psi(P) = P' - P,$$

$$\varphi(P) = \psi(\psi(P)) = (P' - P)' - (P' - P).$$

$\varphi(P)$ heißt das (*erste*) *Residuum von P* .

$\psi(P)$ besteht also genau aus denjenigen Häufungspunkten von P , die außerhalb von P liegen. Offenbar gilt $\psi(P) = \text{cl}(P) - P$.

Die ϕ -Operation tilgt viele unkomplizierte Punktmengen komplett: Ist $A \subseteq \mathbb{R}$ abgeschlossen, so ist bereits $\psi(A) = \emptyset$, also auch $\phi(A) = \emptyset$. Ebenso gilt $\phi(U) = \emptyset$ für alle offenen Mengen U , denn dann ist $\psi(U) = U' - U = \text{cl}(U) \cap (\mathbb{R} - U)$ abgeschlossen, also $\phi(U) = \psi(\psi(U)) = \emptyset$.

Hausdorff (1914): „Nach Definition ist $\psi(M) \subseteq M_\beta [= M']$, also $\phi(M) \subseteq M_{\beta\beta} [= M'']$ und jedenfalls $\phi(M) \subseteq M_h [M \cap M']$; auch dieser Prozeß befreit M also von den isolierten Punkten. Aber er ist im allgemeinen viel energischer als die Kohärenzbildung, in dem er z. B. auch alle inneren Punkte abspaltet [alle $x \in M$ mit $x \in \text{int}(M)$].“

Wenden wir ψ zweimal an, so erhalten wir eine relativ abgeschlossene Teilmenge der Ausgangsmenge. Diese und andere nützliche Eigenschaften versammelt der folgende Satz.

Satz (*Eigenschaften der ψ - und ϕ -Operationen*)

Sei $P \subseteq \mathbb{R}$. Dann gilt:

- (i) $\phi(P) \subseteq P$,
- (ii) $\phi(P)$ ist abgeschlossen in P (d. h. $\phi(P)' \cap P \subseteq \phi(P)$),
- (iii) $\text{cl}(\phi(P)) \subseteq \text{cl}(\psi(P)) \subseteq \text{cl}(P)$,
- (iv) $\psi(\phi(P)) = \phi(\psi(P))$.

Beweis

zu (i):

Ist $x \in \phi(P)$, so ist $x \in (P' - P)' \subseteq P'' \subseteq P'$, also $x \in P'$.
 Andererseits ist $\phi(P) = (P' - P)' - (P' - P)$, also ist $x \notin P' - P$.
 Also gilt $x \in P'$ und $x \notin (P' - P)$, also $x \in P$.

zu (ii):

Es gilt $\phi(P) \subseteq (P' - P)'$, also auch $\phi(P)' \subseteq (P' - P)'$.
 Sei $x \in \phi(P)' \cap P$. Dann gilt $x \in (P' - P)'$.
 Wegen $x \in P$ ist aber $x \notin P' - P$.
 Also $x \in (P' - P)' - (P' - P) = \phi(P)$.

zu (iii):

Wegen $\psi(P) \subseteq \text{cl}(P)$ ist auch $\text{cl}(\psi(P)) \subseteq \text{cl}(P)$.
 Analog ist wegen $\phi(P) = \psi(\psi(P)) \subseteq \text{cl}(\psi(P))$ auch $\text{cl}(\phi(P)) \subseteq \text{cl}(\psi(P))$.

zu (iv):

— Es gilt $\psi(\phi(P)) = \psi(\psi(\psi(P))) = \phi(\psi(P))$.

Keine Überraschung an dieser Stelle ist die Definition der iterierten Residuenbildung eines beliebigen $P \subseteq \mathbb{R}$:

Definition (*iterierte Residuenbildung, $\varphi_\alpha(P)$*)

Sei $P \subseteq \mathbb{R}$. Dann ist das α -te Residuum $\varphi_\alpha(P)$ von P für Ordinalzahlen α rekursiv definiert durch:

$$\begin{aligned}\varphi_0(P) &= P, \\ \varphi_{\alpha+1}(P) &= \varphi(\varphi_\alpha(P)) \quad \text{für } \alpha \text{ Ordinalzahl,} \\ \varphi_\lambda(P) &= \bigcap_{\alpha < \lambda} \varphi_\alpha(P) \quad \text{für } \lambda \text{ Limesordinalzahl.}\end{aligned}$$

Nach dem Satz oben gilt $\varphi_0(P) \supseteq \varphi_1(P) \supseteq \dots \supseteq \varphi_\alpha(P) \supseteq \varphi_{\alpha+1}(P) \supseteq \dots$

Gilt $\varphi_{\alpha+1}(P) = \varphi_\alpha(P)$ für ein α , so gilt $\varphi_\beta(P) = \varphi_\alpha(P)$ für alle $\beta \geq \alpha$. Wieder wird ein endgültiges Ergebnis in abzählbar vielen Schritten erreicht:

Satz (*abzählbare Terminierung der Residuenbildung*)

Sei $P \subseteq \mathbb{R}$. Dann existiert ein abzählbares α mit $\varphi_{\alpha+1}(P) = \varphi_\alpha(P)$.

Beweis

Andernfalls ist $\langle \varphi_\alpha(P) \mid \alpha < \omega_1 \rangle$ eine $\subset$ -absteigende Folge von in P

– abgeschlossenen Teilmengen von P , *Widerspruch*.

Definition (*letztes Residuum, reduzible Mengen*)

Sei $P \subseteq \mathbb{R}$, und sei α derart, daß $\varphi_\alpha(P) = \varphi_{\alpha+1}(P)$ gilt.

Dann heißt $\varphi_\alpha(P)$ das *letzte Residuum* von P .

P heißt *reduzibel*, falls das letzte Residuum von P leer ist.

Reduzible Punktmengen sind also solche, bei denen die „energische“ Residuenbildung irgendwann alles gelöscht hat. Jedes abgeschlossene reduktible P und jedes $P \subseteq \mathbb{R}$ mit verschwindender letzter Kohärenz ist auch reduzibel, aber nicht umgekehrt.

Im Unterschied zur Ableitung oder zur Kohärenzbildung ist die Residuenbildung nicht mehr monoton: Aus $P \subseteq Q$ folgt i.a. nicht, daß $\varphi(P) \subseteq \varphi(Q)$ – z. B. ist $\varphi(Q) = \emptyset$ für $Q = \mathbb{R}$. Zumindest gilt aber eine relative Monotonieeigenschaft.

Satz (*Monotonie der Residuenbildung für relativ abgeschlossene Mengen*)

Sei $P \subseteq \mathbb{R}$. Dann ist $\varphi \mid \mathcal{F}_P : \mathcal{F}_P \rightarrow \mathcal{F}_P$ monoton, d. h. für alle $A, B \in \mathcal{F}_P$ gilt: Ist $A \subseteq B$, so ist $\varphi(A) \subseteq \varphi(B)$.

Beweis

Seien $A, B \in \mathcal{F}_P$, und sei $A \subseteq B$. Dann ist A abgeschlossen in B . Es gilt:

$$\psi(A) = A' - A = A' - B \subseteq B' - B = \psi(B),$$

denn ist $x \in B - A$, so ist $x \notin A'$ wegen A abgeschlossen in B . Analog ist

$$\varphi(A) = \psi(A)' - \psi(A) = \psi(A)' - \psi(B) \subseteq \psi(B)' - \psi(B) = \varphi(B),$$

denn $\psi(A)$ ist abgeschlossen in $\psi(B)$:

Ist $x \in \psi(A)' \cap \psi(B)$, so ist $x \notin B$, also $x \notin \varphi(A) \subseteq A \subseteq B$.

– Also $x \notin \psi(A)' - \psi(A)$, also $x \in \psi(A)$.

Die reduziblen Teilmengen der reellen Zahlen haben eine sehr ansprechende Charakterisierung als Differenzenketten aus abgeschlossenen Mengen:

Satz (*Charakterisierung der reduziblen Mengen*)

Sei $P \subseteq \mathbb{R}$. Dann sind äquivalent:

- (i) P ist reduzibel.
- (ii) Es gibt ein $\gamma < \omega_1$ und eine $\subseteq$ -absteigende Folge $\langle A_\alpha \mid \alpha < 2^\gamma \rangle$ von abgeschlossenen Mengen mit:

$$P = \bigcup_{\alpha < \gamma} A_{2\alpha} - A_{2\alpha+1}.$$

Beweis

(i) $\hookrightarrow$ (ii):

Für alle $Q \subseteq \mathbb{R}$ gilt $\text{cl}(\varphi(Q)) \subseteq \text{cl}(\psi(Q)) \subseteq \text{cl}(Q)$ nach (iii) im Satz oben, und wir haben die Zerlegung:

$$(+)\quad Q = \text{cl}(Q) - \text{cl}(\psi(Q)) \cup \varphi(Q).$$

$$\text{Denn } Q = \text{cl}(Q) - \psi(Q) = \text{cl}(Q) - (\psi(Q) \cup \varphi(Q)) \cup \varphi(Q) = \text{cl}(Q) - \text{cl}(\psi(Q)) \cup \varphi(Q).$$

Sei γ minimal mit $\varphi_\gamma(P) = \emptyset$. Wir setzen für $\alpha < \gamma$:

$$A_{2\alpha} = \text{cl}(\varphi_\alpha(P)),$$

$$A_{2\alpha+1} = \text{cl}(\psi(\varphi_\alpha(P))).$$

Dann ist $\langle A_\alpha \mid \alpha < 2^\gamma \rangle$ eine $\subseteq$ -absteigende Folge von abgeschlossenen Mengen, und aus (+) und $\varphi_\gamma(P) = \emptyset$ folgt

$$P = \bigcup_{\alpha < \gamma} A_{2\alpha} - A_{2\alpha+1}.$$

(ii) $\hookrightarrow$ (i):

Entscheidend ist:

(++) Seien $A, B, C \subseteq \mathbb{R}$ abgeschlossen, $C \subseteq B \subseteq A$, und sei $D \subseteq (A - B) \cup C$ abgeschlossen in $(A - B) \cup C$. Dann gilt $\varphi(D) \subseteq C$.

Beweis von (++)

Es gilt wegen $D \cap (A - B)$ abgeschlossen in $A - B$:

$$\psi(D) = D' - D \subseteq D' - (A - B) \subseteq (D' - A) \cup B = B.$$

Dann aber auch $\psi(D)' \subseteq B$ wegen B abgeschlossen, und somit

$$\varphi(D) = \psi(D)' - \psi(D) \subseteq \psi(D)' \subseteq B.$$

Aber $\varphi(D) \subseteq D$. Also $\varphi(D) \subseteq D \cap B \subseteq C$.

Durch Induktion über $\beta < \gamma$ folgt aber aus (++):

$$\varphi_\beta(P) \subseteq \bigcup_{\beta \leq \alpha < \gamma} A_{2\alpha} - A_{2\alpha+1}$$

Also $\varphi_\gamma(P) \subseteq \bigcap_{\alpha < 2^\gamma} A_\alpha$, eine abgeschlossene Menge.

Also gilt $\varphi_{\gamma+1}(P) \subseteq \varphi(\bigcap_{\alpha < 2^\gamma} A_\alpha) = \emptyset$,

ersteres wegen $\varphi_\gamma(P)$ abgeschlossen in $\bigcap_{\alpha < 2^\gamma} A_\alpha$,

— letzteres wegen $\varphi(A) = \emptyset$ für alle abgeschlossenen Mengen A .

Es gibt eine weitere schöne Charakterisierung der reduziablen Mengen, die wir hier ohne Beweis angeben (vgl. Hausdorff 1914, 1927). Sei $P \subseteq \mathbb{R}$. Dann sind äquivalent:

- (i) P ist reduziabel.
- (ii) $P \in \Delta_2 = \Sigma_2 \cap \Pi_2 = \mathcal{F}_\sigma \cap \mathcal{G}_\delta$ (d.h. P ist sowohl eine abzählbare Vereinigung von abgeschlossenen Mengen, als auch ein abzählbarer Schnitt von offenen Mengen.)

Der Beweis dieses Satzes ist alles andere als schwer, benötigt aber ein wenig mehr Begrifflichkeit, als wir hier entwickelt haben. Er zeigt die Natürlichkeit der Residuenbildung. Wir erhalten durch Betrachtung der kleinsten Stelle γ mit $\varphi_\gamma(P) = \emptyset$ für $P \in \Delta_2$ eine sehr feine Hierarchie der Δ_2 -Mengen, und gewinnen so ein Gefühl für die Komplexität der Borelmengen: Schon der Übergang von den offenen und den abgeschlossenen Mengen zur nächsten Komplexitätsstufe Δ_2 birgt einen unerwarteten strukturellen Reichtum. Und die Borel-Hierarchie hat ω_1 -viele Stufen...

Felix Hausdorff über Punktmengen und Ordinalzahlen

„Bei G. Cantor spielt in der Lehre von den Punktmengen die Theorie der Ordnungszahlen eine wesentliche Rolle, ja sie ist wohl ursprünglich eben zu diesem Zweck ausgebildet worden. E. Lindelöf erkannte, daß der Begriff der Verdichtungspunkte [Kondensationspunkte] (der übrigens schon bei Cantor [1885d] vorkommt), die Ordnungszahlen [für einen Beweis des Satzes von Cantor-Bendixson] auszuschalten gestattet, und dieser Tendenz folgt auch unsere Darstellung, indem sie den Kern A_k , der bei Cantor als Endglied einer wohlgeordneten Reihe von Mengen erscheint, mit einem Schlage als größte insichdichte Teilmenge von A definiert. Dennoch werden die Ordnungszahlen ihren Platz in der Punktmengentheorie behaupten, aus historischen und sachlichen Gründen. Sie ermöglichen eine feinere Analyse der Struktur gegebener Mengen durch Unterscheidung von Häufungspunkten verschieden hoher Ordnung; eine Analyse, die z. B. bei funktionentheoretischen Untersuchungen wertvolle Dienste geleistet hat (wie bei den Mittag-Lefflerschen Existenzbeweisen für Funktionen mit vorgeschriebenen singulären Stellen), freilich aber dem Mächtigkeitsproblem der Punktmengen [dem Kontinuumsproblem] keinen Schritt näher kommt als die andere, mehr summarische Behandlung des Unabzählbaren ...“

(Felix Hausdorff 1914, „Grundzüge der Mengenlehre“)

12. Die Mächtigkeiten abgeschlossener Mengen

In diesem Kapitel untersuchen wir die perfekten und allgemeiner die abgeschlossenen Mengen genauer. Wir wissen, daß nichtleere perfekte Mengen überabzählbar sind. Nun zeigen wir stärker, daß sie immer die Mächtigkeit des Kontinuums besitzen. Aus der Cantor-Bendixson-Zerlegung können wir dann die möglichen Mächtigkeiten der abgeschlossenen Mengen ablesen.

Jede nichtleere offene Menge enthält ein reelles Intervall. (Hier und im folgenden meint „P enthält ein Intervall“ immer „P enthält ein nichttriviales Intervall“, d. h. es gibt $a, b \in \mathbb{R}$ mit $a < b$ und $[a, b] \subseteq P$. Dies vereinfacht viele Formulierungen.) Gilt etwas Ähnliches für die abgeschlossenen Mengen? Beispiele für abgeschlossene Mengen sind abzählbare Punktmengen wie $\{0\}$, $\{1/n \mid n \geq 1\} \cup \{0\}$, $\mathbb{Z}$ usw., und des weiteren abgeschlossene Intervalle $[a, b]$. Sind nun alle abgeschlossenen Mengen Kombinationen dieser beiden Typen? Konkreter: Enthält jede überabzählbare abgeschlossene Menge ein Intervall? Es ist ein nichttriviales Problem, ein Gegenbeispiel zu konstruieren! Ein Gegenbeispiel werden wir nun kennenlernen, und später zeigen, daß es *das* Gegenbeispiel ist.

Die Cantormenge

Der Beweis des Satzes über die Mächtigkeit der perfekten Mengen benutzt die Idee der Cantormenge C . Diese Menge ist uns bereits bei der Diskussion von $|\mathbb{R}| = |\mathbb{N}|^2 = |\mathcal{P}(\mathbb{N})|$ begegnet (1.9). Hier noch einmal die Definition.

Definition (*Cantormenge*)

Die *Cantormenge* C besteht aus allen reellen Zahlen x , die eine (nicht notwendig kanonische) Ternärdarstellung $x = 0, c_1 c_2 c_3 \dots$ besitzen mit $c_n \in \{0, 2\}$ für alle $n \geq 1$.
 C heißt auch das *Cantorsche Diskontinuum*.

Offenbar gilt $C \subseteq [0, 1]$.

Für eine reelle Zahl $x = 0, d_1 d_2 d_3 \dots$ in Dezimaldarstellung gilt
$$x = \sup \{ d_1/10 + d_2/10^2 + \dots + d_n/10^n \mid n \geq 1 \}.$$

Ebenso gilt für eine reelle Zahl $x = 0, c_1 c_2 c_3 \dots$ in Ternärdarstellung (d. h. 3-adischer Darstellung), daß

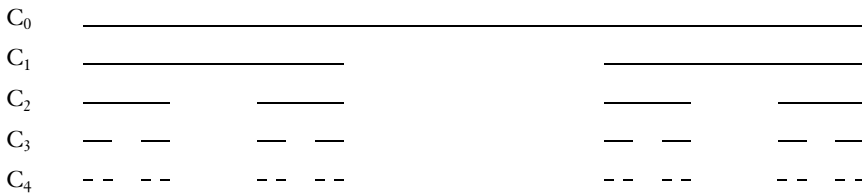
$$x = \sup \{ c_1/3 + c_2/3^2 + \dots + c_n/3^n \mid n \geq 1 \}.$$

Die Cantormenge ist dann die Menge all derer Summen, in denen c_n nur die Werte 0 oder 2 annimmt, was wir suggestiv wie folgt schreiben können:

$$C = \{x \in \mathbb{R} \mid x = c_1/3 + c_2/3^2 + \dots + c_n/3^n + \dots, \text{ mit } c_n \in \{0, 2\} \text{ für alle } n \geq 1\}.$$

Der mittlere der möglichen Werte 0, 1 und 2 wird in den Ternärdarstellungen von $x \in C$ gerade vermieden. Diese arithmetische Eigenschaft hat ein geometrisches Gegenstück: Ist die erste Nachkommastelle eines $x \in C$ gleich 0, so liegt x im ersten Drittel $[0, 1/3]$ des Einheitsintervalls; ist sie gleich 2, so liegt x im dritten Drittel $[2/3, 1]$. Ist die erste und die zweite Nachkommastelle gleich 0, so liegt x im ersten Drittel von $[0, 1/3]$, also gilt $x \in [0, 1/9]$. Ist die erste Nachkommastelle 0, die zweite 2, so liegt x im dritten Drittel von $[0, 1/3]$, also gilt $x \in [2/9, 1/3]$, usw. usf.

Die Cantormenge kann man auf diese Weise durch iteriertes Entfernen von mittleren Dritteln aus Intervallketten anschaulich machen – und dabei zugleich Sympathie für diese Menge entwickeln. Wir starten mit dem Einheitsintervall $[0, 1]$ und entfernen dann wiederholt die offenen mittleren Drittelintervalle aus $[0, 1]$ und den verbleibenden Resten:



Es gilt also für alle $n \in \mathbb{N}$:

$$C_{n+1} = C_n \cap \bigcup_{0 \leq k < 3^{n+1}, k \text{ gerade}} [k/3^{n+1}, (k+1)/3^{n+1}],$$

wobei $C_0 = [0, 1]$.

Alle Teilintervalle der Mengen C_n sind abgeschlossen, und $\langle C_n \mid n \in \mathbb{N} \rangle$ ist damit eine $\subseteq$ -absteigende Folge von abgeschlossenen, nichtleeren und beschränkten Mengen. Also ist auch ihr Schnitt

$$C_\omega = \bigcap_{n \in \mathbb{N}} C_n$$

abgeschlossen und nichtleer. C_ω ist nun gerade die Cantormenge:

Übung

Es gilt $C_\omega = C$.

Definition (natürliche Approximation von C)

Die Folge $\langle C_n \mid n \in \mathbb{N} \rangle$ heißt die *natürliche Approximation der Cantormenge C* .

Bemerkenswert ist weiter, daß die Cantormenge die Ableitung der linken (oder auch der rechten) Randpunkte der Intervalle in den Mengen C_n ist:

Übung

Sei $P = \{x \in [0, 1] \mid x = 0, c_1 \dots c_n \text{ in Ternärdarstellung}$
mit $c_i \in \{0, 2\}$ für $1 \leq i \leq n, n \in \mathbb{N}\}$.

Dann ist $P \subseteq \mathbb{Q}$ und es gilt $P' = C$.

Die Menge C hat Cantor in einer Anmerkung des fünften Teils der „Linearen Punktmannigfaltigkeiten“ eingeführt (1883b, Anmerkung 11), als Beleg für seine Behauptung, daß eine nichtleere perfekte Menge keineswegs ein Intervall enthalten muß.

Cantor (1883b): „Als ein Beispiel einer perfekten Punktmenge, die in keinem noch so kleinen Intervall überall dicht ist, führe ich den Inbegriff aller reellen Zahlen an, die in der Formel:

$$z = c_1/3 + c_2/3^2 + \dots + c_v/3^v + \dots$$

enthalten sind, wo die Koeffizienten c_v nach Belieben die beiden Werte 0 und 2 annehmen haben und die Reihe sowohl aus einer endlichen, wie aus einer unendlichen Anzahl von Gliedern bestehen kann.“

Die Behauptungen „perfekt“ und „in keinem Intervall dicht“ zeigt der folgende Satz.

Satz (*Eigenschaften der Cantormenge C*)

- (i) C ist perfekt.
- (ii) C enthält kein Intervall $[a, b]$ für $a, b \in \mathbb{R}$ mit $a < b$.
- (iii) Es gilt $|C| = |\mathbb{R}|$.

Beweis

zu (i): Wegen $C = C_\omega = \bigcap_{n \in \omega} C_n$ ist C abgeschlossen.

(Dies kann man auch direkt aus der Definition beweisen:

Ein Häufungspunkt von C hat eine Ternärdarstellung, in der keine 1 als Nachkommastelle vorkommt.)

Wir müssen noch zeigen, daß jedes $x \in C$ ein Häufungspunkt von C ist.

Sei also $x \in C$, und sei

$x = 0, c_1 c_2 c_3 \dots$ mit $c_n \in \{0, 2\}$ für alle $n \geq 1$.

Ist $\{n \in \mathbb{N} \mid c_n \neq 0\}$ unbeschränkt in $\mathbb{N}$, so sei $x_n = 0, c_1 \dots c_n$ für $n \geq 1$.

Dann ist $X = \{x_n \mid n \geq 1\} \subseteq C$, $x \notin X$ und $x = \sup(X)$.

Also ist x Häufungspunkt von C .

Andernfalls sei $k = \min \{n \geq 1 \mid c_m = 0 \text{ für alle } m \geq n\}$.

Weiter setzen wir $y = 0, c_1 \dots c_{k-1}$ (also $y = 0$ für $k = 1$).

Für $n \geq k$ sei $x_n = y + 2/3^n$, also $x_n = 0, c_1 \dots c_{k-1} 0 \dots 0 2$, wobei die letzte 2 die n -te Nachkommastelle ist.

Dann ist $X = \{x_n \mid n \geq k\} \subseteq C$, $x \notin X$ und $x = \inf(X)$.

Also ist x Häufungspunkt von C .

zu (ii): Sei $[a, b] \subseteq [0, 1]$, $a < b$, und sei $\varepsilon = b - a$.

Sei $n \in \mathbb{N}$ mit $1/3^n < \varepsilon/2$.

Dann existiert ein $k \in \mathbb{N}$ mit $0 \leq k < 3^n$ und

$$[k/3^n, (k+1)/3^n] \subseteq [a, b].$$

Sei $I = [k/3^n, (k+1)/3^n]$.

Das offene mittlere Drittelintervall $] (3k+1)/3^{n+1}, (3k+2)/3^{n+1} [$ von I hat aber mit $C = C_\omega$ leeren Schnitt, also kann wegen $I \subseteq [a, b]$ nicht $[a, b] \subseteq C$ gelten.

zu (iii): Jedes $x \in C$ hat eine eindeutige Darstellung $x = 0, c_1 c_2 c_3 \dots$ mit $c_n \in \{0, 2\}$ für $n \geq 1$ (!), und jedes so darstellbare x ist in C .

— Also gilt $|C| = |\mathbb{N} - \{0\}| \{0, 2\}^\mathbb{N} = 2^\omega = |\mathbb{R}|$.

Die Cantormenge ist damit ein Beispiel für eine überabzählbare abgeschlossene Menge, die kein Intervall enthält. C ist ein Zwitterwesen: Es ist weder aus verborgen isolierten Punkten noch aus Abschnitten der Linie zusammengesetzt. Derartige Mengen lassen sich durch die Methode der Entfernung von Intervallen vielfältig erzeugen. Umgekehrt werden wir unten zeigen, daß alle nichtleeren perfekten Mengen, die keine Intervalle enthalten, mit dieser Methode erzeugt werden können.

Aufgrund der Wichtigkeit der Cantormenge vergeben wir ein festes Zeichen für ihren Ordnungstyp unter der von den reellen Zahlen ererbten Ordnung.

Definition (der Ordnungstyp der Cantormenge)

Wir setzen:

$$\xi = \text{o.t.}(\langle C, < \rangle).$$

Übung

Es gilt $\xi = \exp_{\text{lex}}(2, \omega)$.

Nirgendso dichte und magere Mengen

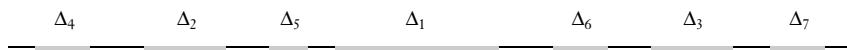
Die Frage, ob jede überabzählbare abgeschlossene Menge ein Intervall enthält, wird durch die Cantormenge also negativ beantwortet. Eine Abschwächung der Frage ist: *Enthält eine überabzählbare abgeschlossene Menge immer eine rationale Zahl als Element?* Für die Cantormenge ist dies richtig. Im allgemeinen ist jedoch auch diese Eigenschaft nicht erfüllt:

Hausdorff (1914): „Aber auch abgeschlossene Mengen mit nichtverschwindendem Kern, z. B. perfekte Mengen, können nirgendso dicht sein [müssen keine nichttrivialen Intervalle enthalten]. Um auf der geraden Linie eine solche zu bilden, verstehen wir unter Δ eine offene Strecke ($a < x < b$) und unter

$$G = \Delta_1 + \Delta_2 + \dots = \sum \Delta_n$$

eine Summe von abzählbar vielen, paarweise fremden Strecken, also ein lineares Ge-

biet [eine offene Teilmenge von $\mathbb{R}$] ... Damit das Komplement F perfekt werde, vermeiden wir bei der sukzessiven Wahl der Strecken $\Delta_1, \Delta_2, \dots$, daß zwei von ihnen mit einem Endpunkt zusammenstoßen, der ja sonst ein isolierter Punkt von F sein würde;



die Figur zeigt den Anfang einer solchen Streckenwahl. Um endlich G dicht, also F nirgendsdicht zu machen, lassen wir G etwa die Menge $R = \{r_1, r_2, \dots\}$ der rationalen Zahlen einschließen; wir wählen alle Δ_n mit irrationalen Endpunkten, bedecken r_1 mit Δ_1 , das erste (mit niedrigstem Index behaftete) von Δ_1 nicht eingeschlossene r_p mit Δ_2 , das erste von $\Delta_1 + \Delta_2$ nicht eingeschlossene r_γ mit Δ_3 usw. Hiermit wird $R \subseteq G$, also G auf der ganzen Linie dicht. Wir werden sehen, daß die perfekte Menge F die Mächtigkeit des Kontinuums hat. Mengen so hoher Mächtigkeit können also nirgendsdicht sein. Man kann sogar, wenn δ_n die Länge der Strecke Δ_n bedeutet, die Summe $\sum \delta_n$ (das ‚Längenmaß‘ von G) konvergent und beliebig klein machen, was ja mit der angegebenen Konstruktion verträglich ist; Gebiete [offene Mengen] von beliebig kleinem Längenmaß können also die ganze Gerade dicht erfüllen.

Natürlich kann man diese Konstruktion mannigfach abändern. Ein Beispiel (wo allerdings $\sum \delta_n$ divergiert) ist das folgende: die Menge bestehe aus allen Zahlen, die eine Dezimalbruchentwicklung ... zulassen, in der eine der Ziffern 1 – 8 nicht vorkommt, etwa die Ziffer 3 ... Das erste Beispiel dieser Art wurde von $G.$ Cantor gegeben, nämlich die Menge der Zahlen, in deren triadischer Entwicklung ... die Ziffer 1 nicht vorkommt.“

Die hier rekursiv konstruierte Menge $F = \mathbb{R} - G = \mathbb{R} - \bigcup_{n \in \mathbb{N}} \Delta_n$ ist also eine perfekte Menge mit $F \cap \mathbb{Q} = \emptyset$. Während der Konstruktion von F wird sichergestellt, daß alle rationalen Zahlen in einem entfernten Intervall (einem Δ_n) liegen.

Die angedeutete Variation des Längenmaßes von G während dieser Konstruktion wollen wir noch etwas genauer betrachten. Sei $\sigma \in \mathbb{R}$, $\sigma > 0$. Weiter wählen wir eine Folge $\langle \varepsilon_n \mid n \in \mathbb{N} \rangle$ von positiven reellen Zahlen, deren Summe gegen σ konvergiert, etwa $\varepsilon_n = \sigma/2^{n+1}$ für $n \in \mathbb{N}$. Bei der von Hausdorff durchgeführten Konstruktion können wir immer $\Delta_n =]a_n, b_n[$ so klein wählen, daß das Maß $b_n - a_n$ von Δ_n kleinergleich ε_n wird. Dann haben wir während der Konstruktion von $F = \mathbb{R} - \bigcup_{n \in \mathbb{N}} \Delta_n$ insgesamt höchstens eine Menge vom Maß σ aus $\mathbb{R}$ entfernt. F ist also fast ganz $\mathbb{R}$, enthält aber dennoch kein Intervall. Umgekehrt ist das Komplement $G = \bigcup_{n \in \mathbb{N}} \Delta_n$ allenfalls von der Länge σ , enthält aber für jedes nichtleere reelle Intervall I eine Menge $\Delta_n \cap I \neq \emptyset$, also ein Teilintervall von I !

Die Diskussion legt folgende Begriffe nahe.

Definition (*nirgendsdicht, mager*)

Sei $P \subseteq \mathbb{R}$. P heißt *nirgendsdicht*, falls $\text{cl}(P) = P \cup P'$ keine nichttrivialen Intervalle enthält.

P heißt *mager*, falls P eine abzählbare Vereinigung von nirgendsdichten Mengen ist. P heißt *komager*, falls $\mathbb{R} - P$ mager ist.

P nirgends dicht ist äquivalent zu $\text{int}(\text{cl}(P)) = \emptyset$. P ist nirgends dicht genau dann, wenn $\text{cl}(P)$ nirgends dicht ist. Die Cantormenge C ist nirgends dicht.

Magere Mengen können dicht in $\mathbb{R}$ sein: Jede abzählbare Menge ist mager, und insbesondere ist also $\mathbb{Q}$ mager. Dagegen zeigt eine einfache Variante von Cantors erstem Beweis der Überabzählbarkeit von $\mathbb{R}$, daß die Bezeichnung „mager“ wirklich gerechtfertigt ist: $\mathbb{R}$ ist nicht mager. Die mageren Mengen bilden dann offenbar ein ω_1 -vollständiges Ideal auf $\mathbb{R}$ (und die komageren Mengen einen ω_1 -vollständigen Filter).

Satz (*Bairescher Kategoriensatz*)

Sei $U \subseteq \mathbb{R}$ offen und nichtleer. Dann ist U nicht mager.

Beweis

Sei $U \subseteq \mathbb{R}$ offen und nichtleer, und seien $P_n \subseteq \mathbb{R}$ nirgends dicht für $n \in \mathbb{N}$.

Wir konstruieren ein $x^* \in U$ mit $x^* \notin P_n$ für alle $n \in \mathbb{N}$.

Dann ist also $\bigcup_{n < \omega} P_n \neq U$. U ist also nicht mager.

Sei $I_k =]a_k, b_k[$, $k < \omega$, eine Aufzählung aller nichttrivialen rationalen Intervalle. Wir definieren rekursiv k_n für $n < \omega$ durch:

$k_0 =$ „das kleinste k mit $[a_k, b_k] \subseteq U$ “,

$k_{n+1} =$ „das kleinste $k > k_n$ mit $[a_k, b_k] \subseteq I_{k_n} - P_n$ “.

[Ein solches k existiert, da P_n nicht dicht in I_{k_n} ist.]

Sei $[x, y] = \bigcap_{k < \omega} [a_{k_n}, b_{k_n}]$, also $x = \sup_{n < \omega} a_{k_n}$, $y = \inf_{n < \omega} b_{k_n}$.

— Dann ist $x \leq y$ und $[x, y] \cap \bigcup_{n < \omega} P_n = \emptyset$. Also ist $x^* = x$ wie gewünscht.

Der Bairesche Kategoriensatz findet sich in [Baire 1899]. Magere Mengen nennt man auch „von erster Kategorie“, nicht magere Mengen „von zweiter Kategorie“. Der Satz sagt dann: Nichttriviale offene Mengen sind von zweiter Kategorie.

Der Beweis ist lediglich eine Verfeinerung des ersten Cantorschen Beweises der Überabzählbarkeit von $\mathbb{R}$. Wir erhalten also keinen wirklich neuen Beweis der Überabzählbarkeit von $\mathbb{R}$ aus dem Baireschen Kategoriensatz durch das Argument: Jede abzählbare Menge ist mager. $\mathbb{R}$ ist jedoch nicht mager, also überabzählbar. Andererseits suggeriert der Beweis die folgende Version von Cantors erstem Beweis der Überabzählbarkeit der reellen Zahlen:

Variante des ersten Cantorschen Beweises von „ $\mathbb{R}$ ist überabzählbar“

Für reelle $a < b$ seien $L(I) = [a, a + (b - a)/3]$ und $R(I) = [b - (b - a)/3, b]$ das linke und rechte Drittelintervall von $I = [a, b]$. (De facto sind zwei beliebige disjunkte abgeschlossene nichtleere Teilintervalle von I geeignet.)

Seien nun $x_0, x_1, \dots, x_n, \dots$, $n < \omega$, reelle Zahlen in $I_0 = [a, b]$.

Wir definieren rekursiv Intervalle I_n für $n < \omega$ durch:

$$I_{n+1} = \begin{cases} L(I_n), & \text{falls } x_n \in R(I_n), \\ R(I_n), & \text{sonst.} \end{cases}$$

Sei $J = \bigcap_{n < \omega} I_n$. Dann ist $J \neq \emptyset$ und $x_n \notin J$ für alle $n < \omega$.

Also ist I überabzählbar.

Abgeschlossene Teilmengen von $\mathbb{R}$ sind genau dann nicht mager, wenn sie ein nichttriviales Intervall enthalten. Andernfalls sind sie sogar nirgends dicht.

Für den dualen Begriff der komageren Mengen gilt:

Übung

(a) Sei $A \subseteq \mathbb{R}$. Dann sind äquivalent:

(i) A ist nirgends dicht.

(ii) $\mathbb{R} - A$ enthält eine offene dichte Teilmenge.

[direkt oder mit: $\text{int}(\mathbb{R} - A) = \mathbb{R} - \text{cl}(A)$, $\text{cl}(\mathbb{R} - A) = \mathbb{R} - \text{int}(A)$.]

(b) Ist $A \subseteq \mathbb{R}$ komager, so ist A dicht in $\mathbb{R}$.

(c) Ist U_n offen und dicht in $\mathbb{R}$ für $n \in \mathbb{N}$, so ist

$\bigcap_{n \in \mathbb{N}} U_n$ komager, und insbesondere dicht in $\mathbb{R}$.

Der Schnitt über abzählbar viele offene dichte Teilmengen von $\mathbb{R}$ ist i. a. nicht mehr offen, und kann sogar ein leeres Inneres haben: Betrachte etwa $U_q = \mathbb{R} - \{q\}$ für $q \in \mathbb{Q}$. Er ist aber immer komager, was viel stärker ist als dicht.

Obige Konstruktion von Hausdorff zeigt: Es gibt eine Zerlegung von $\mathbb{R}$ in Mengen A, B mit: A hat Maß Null, B ist mager. Hierzu wird die Konstruktion abzählbar oft durchgeführt, mit einer Folge von $\sigma_n > 0$, die gegen Null konvergiert. Diese liefert jeweils nirgends dichte F_n und offene dichte G_n wie oben. Dann hat $A = \bigcap_{n < \omega} G_n$ ein Maß kleiner σ_n für alle n , also Maß 0, und $B = \mathbb{R} - A = \bigcup_{n < \omega} F_n$ ist mager. Die beiden Begriffe einer kleinen Teilmenge von $\mathbb{R}$, die durch „Maß Null“ und „mager“ gegeben werden – d. h. mathematisch: die zugehörigen Ideale – sind also völlig verschieden.

Wir haben hier den Begriff „Längenmaß“ eher intuitiv behandelt. Eine kurze offizielle Definition von „ A hat (Lebesguesches) Längenmaß 0“ für ein $A \subseteq \mathbb{R}$ ist: Für alle $\sigma \in \mathbb{R}$, $\sigma > 0$ gibt es Intervalle $I_n = [a_n, b_n]$, $n < \omega$, $a_n \leq b_n$, mit $A \subseteq \bigcup_{n < \omega} I_n$ und $\sum_{n < \omega} b_n - a_n < \sigma$. Die Mengen vom Maß 0 bilden dann ein ω_1 -vollständiges Ideal, wie man leicht zeigt.

Bevor wir nun den Ordnungstyp ξ der Cantormenge genauer untersuchen, wollen wir mit dem durch die Cantormenge geschärften Auge noch einmal einen Blick auf die verborgen isolierten Punkte einer abgeschlossenen Menge werfen – und auf die Fallstricke der Theorie.

Die Ableitung der Restmenge

Ein reduktibles P ist immer abzählbar. Das Beispiel $P = \mathbb{Q}$ dagegen zeigt, daß nicht jedes abzählbare P durch iteriertes Ableiten zum Verschwinden gebracht werden kann. Dagegen scheint für jedes abgeschlossene $P \subseteq \mathbb{R}$ die Menge der verborgen isolierten Punkte von P ein guter Kandidat für eine reduktible Menge zu sein:

Stellen wir P wie in der Cantor-Bendixson-Zerlegung als $P = P^* \cup A$ dar, so läßt sich vermuten, daß A als eine abzählbare Menge, die durch Aufsammeln der isolierten Punkte der Ableitungen $P^{(\alpha)}$ von P gewonnen wird, selbst reduktibel ist.

Bei Cantor liest man im fünften Teil der Reihe „Über unendliche lineare Punktmannigfaltigkeiten“ ohne Beweis folgende Ankündigung:

Cantor (1883b): „Hat aber $P^{(1)}$ die Mächtigkeit der zweiten Zahlenklasse [Ist $P^{(1)}$ überabzählbar], so läßt sich $P^{(1)}$ stets und zwar nur auf einzige Weise in zwei Mengen R und S zerlegen, so daß:

$$P^{(1)} = R + S \quad [P^{(1)} = R \cup S \text{ und } R \cap S = \emptyset],$$

wo R und S eine äußerst verschiedene Beschaffenheit haben:

R ist so beschaffen, daß sie durch den wiederholten Ableitungsprozeß einer fortwährenden Reduktion bis zur Annihilation fähig ist, so daß es immer eine erste ganze Zahl γ der Zahlenklassen (I) oder (II) gibt [eine abzählbare Ordinalzahl γ], für welche:

$$R^{(\gamma)} = 0 [= \emptyset];$$

solche Punktmengen R nenne ich *reduktibel*.

S ist dagegen so beschaffen, daß bei dieser Punktmenge der Ableitungsprozeß gar keine Änderung hervorbringt, indem:

$$S = S^{(1)}$$

und folglich auch:

$$S = S^{(\gamma)}$$

ist; derartige Mengen S nenne ich *perfekte* Punktmengen.“

Man muß ein wenig suchen, um in dieser Passage eine nach unseren Sätzen nicht bewiesene Aussage zu finden. Aber tatsächlich enthält die zitierte Stelle eine falsche Behauptung: Die Menge A [bei Cantor: R] der verborgen isolierten Punkte von $P' = P^* \cup A$ braucht nicht reduktibel zu sein! Obige Vermutung ist also falsch. Dies ist einem Studenten von Mittag-Leffler, nämlich Ivar Bendixson (1861 – 1935), zuerst aufgefallen. Bereits im nächsten Teil der „Punktmannigfaltigkeiten“ berichtigt Cantor die Behauptung:

Cantor (1884b): „Diese Sätze A, B, C, D, E, F sowohl, wie die hier entwickelten Beweise derselben waren mir zur Zeit der Abfassung von Nr. 5 dieser Abhandlung bekannt; indessen bin ich dort bei der Formulierung des [Zerlegungs-] Satzes E, auf pag. 575, Bd. XXI, etwas zu weit gegangen; so wie der Satz E dort steht, ist er nicht allgemein richtig ...

Diese *wichtige* Bemerkung ist zuerst von Herrn Ivar Bendixson in Stockholm in einem an mich gerichteten Schreiben (v. Mai 1883) gemacht worden ...“

Ein Beispiel, in welchem die abzählbare Restmenge A in der Cantor-Bendixson-Zerlegung $P = P^* \cup A$ einer abgeschlossenen Menge P nicht reduktibel ist, läßt sich nun mit Hilfe der Cantormenge konstruieren. In diesem Beispiel ist sogar A die Menge der isolierten Punkte von P : P^* wird in einem Schritt erreicht.

Satz *(eine perfekte Menge als Ableitung einer Menge aus isolierten Punkten)*

Es existiert eine überabzählbare abgeschlossene Menge $P \subseteq \mathbb{R}$ mit:

Ist $P = P^* \cup A$ die Cantor-Bendixson-Zerlegung von P , so gilt

$$P' = A' = P^*.$$

Beweis

Sei hierzu $\langle C_n \mid n \in \mathbb{N} \rangle$ die natürliche Approximation von C . Weiter sei L die Menge der linken Randpunkte aller Intervalle der Mengen C_n .

Dann ist L abzählbar und $L' = C$ (vgl. obige Übung).

Sei $x^0, x^1, \dots$ eine Aufzählung von L ohne Wiederholungen.

Für jedes $n \in \mathbb{N}$ sei nun $\langle x_k^n \mid k \in \mathbb{N} \rangle$ eine streng monoton wachsende Folge mit den Eigenschaften:

$$(i) \sup_{k \in \mathbb{N}} x_k^n = x^n,$$

$$(ii) \{x_k^n \mid k \in \mathbb{N}\} \subseteq]\sup(C \cap [0, x^n[), x^n[, \text{ falls } x^n \neq 0.$$

Anders ausgedrückt: Die Folge der x_k^n konvergiert von links gegen x^n , und verläuft (für $x^n \neq 0$) ganz in einem mittleren Drittelintervall, das während der Konstruktion von $\langle C_m \mid m \in \mathbb{N} \rangle$ entfernt wurde.

Sei schließlich $A = \{x_k^n \mid n, k \in \mathbb{N}\}$, und

$$P = C \cup A.$$

Dann gilt $P' = C$. Also ist $P = C \cup A$ die Cantor-Bendixson-Zerlegung von P , und insbesondere gilt $A' \subseteq P' = C$.

Andererseits gilt nach Konstruktion $L \subseteq A'$, also

$$C = L' \subseteq A' \text{ wegen } A' \text{ abgeschlossen.}$$

— Also gilt $A' = C$ und damit ist P wie gewünscht.

Für die Menge der verborgen isolierten Punkte einer abgeschlossenen Menge gilt aber, daß ihre Ableitungen irgendwann ganz außerhalb der Menge liegen [Bendixson 1883]:

Übung

Sei $P \subseteq \mathbb{R}$ abgeschlossen und sei $P = P^* \cup A$

die Cantor-Bendixson-Zerlegung von P .

Dann existiert ein abzählbares α mit $A^{(\alpha)} \cap A = \emptyset$.

Die Struktur der perfekten Mengen

Wir zeigen nun, daß nichtleere perfekte Mengen immer die Mächtigkeit des Kontinuums besitzen. Genauer werden wir beweisen, daß eine perfekte Menge, die keine Intervalle enthält, im wesentlichen nichts anderes ist als eine verzerrte Cantormenge.

Hierzu zunächst eine für sich interessante Vorüberlegung:

Definition (*Intervallzerlegung einer offenen Menge*)

Sei $U \subseteq \mathbb{R}$ offen und nichtleer.

Eine Menge $\mathcal{I}$ heißt eine *Intervallzerlegung* von U , falls gilt:

- (i) Jedes $I \in \mathcal{I}$ ist ein offenes, nichtleeres reelles Intervall.
- (ii) $U = \bigcup \mathcal{I}$.
- (iii) $I \cap J = \emptyset$ für alle $I, J \in \mathcal{I}$ mit $I \neq J$.

Satz (*Existenz und Eindeutigkeit der Intervallzerlegung*)

$U \subseteq \mathbb{R}$ offen und nichtleer.

Dann existiert eine eindeutige Intervallzerlegung $\mathcal{I}$ von U .

Beweis

Die Idee ist, von jedem $x \in U$ soweit nach links und rechts zu gehen, bis wir Elemente aus dem Komplement von U treffen. Alle so erhaltenen Intervalle bilden die gesuchte Zerlegung von U .

Sei $A = \mathbb{R} - U$. Für $x \in U$ setzen wir

$$\begin{aligned} \ell_x &= \begin{cases} \sup(\{y \in A \mid y < x\}), & \text{falls } \{y \in A \mid y < x\} \neq \emptyset, \\ -\infty & \text{sonst.} \end{cases} \\ r_x &= \begin{cases} \inf(\{y \in A \mid x < y\}), & \text{falls } \{y \in A \mid x < y\} \neq \emptyset, \\ +\infty & \text{sonst.} \end{cases} \end{aligned}$$

$$I_x =]\ell_x, r_x[.$$

Dann gilt $I_x \subseteq U$ für alle $x \in U$.

Sei nun $\mathcal{I} = \{I_x \mid x \in U\}$. Dann ist $\mathcal{I}$ eine Intervallzerlegung von U .

(i) und (ii) sind klar, und (iii) folgt daraus, daß alle Intervallgrenzen ℓ_x, r_x keine Elemente von U sind. Die Intervalle in $\mathcal{I}$ können sich also nicht überlappen.

— Die Eindeutigkeit einer Intervallzerlegung ist klar.

Eine notationell etwas elegantere Definition von $\mathcal{I}$ ist die folgende.

Für $a, b \in U$ setzen wir:

$a \sim b$ falls $[a, b] \subseteq U$ und $[b, a] \subseteq U$.

Dann ist $\sim$ eine Äquivalenzrelation auf U . Weiter ist jede Äquivalenzklasse $a/\sim$ ein offenes, nichtleeres Intervall, und je zwei verschiedene $a/\sim$ und $b/\sim$ sind disjunkt. Wir setzen $\mathcal{I} = U/\sim = \{a/\sim \mid a \in U\}$. Dann ist $\mathcal{I}$ die gesuchte Intervallzerlegung von U .

Einige Bemerkungen zur Intervallzerlegung $\mathcal{I}$ einer offenen Menge sind:

1. Als Menge von paarweise disjunkten offenen Mengen ist $\mathcal{I}$ abzählbar (jedes $I \in \mathcal{I}$ enthält eine rationale Zahl als Element).

2. Ist $\mathcal{I}$ die Intervallzerlegung von U und ist $I =]c, d[\in \mathcal{I}$, so gilt $c, d \in \mathbb{R} - U$, falls $c, d \in \mathbb{R}$, d. h. $c, d \neq \pm\infty$. Reelle Grenzen der Intervalle in $\mathcal{I}$ gehören also immer zum Komplement von U .

Der zweite Punkt zeigt den Grund, warum ein analoger Satz schon im $\mathbb{R}^2$ falsch ist. Ein $U \subseteq \mathbb{R}^2$ ist *offen*, falls für jedes $x \in U$ ein $\varepsilon > 0$ existiert, sodaß das Innere eines Kreises um x mit Radius ε ganz in U liegt. Ein offenes $U \subseteq \mathbb{R}^2$ läßt sich nun i. a. nicht in *disjunkte* offene Kreise (= Kreisscheiben ohne Rand) zerlegen. Ein Gegenbeispiel ist $U =]0, 1[\times]0, 1[$: Bilden wir für $x \in U$ einen x enthaltenden möglichst großen Kreis K , dessen Inneres ganz in U liegt, so liegen auf der Kreislinie sowohl Punkte von U als auch Punkte von $\mathbb{R} - U$; jedes $y \in U$, das auf der Kreislinie liegt, kann dann nur noch mit Hilfe eines K überlappenden Kreisinneren eingefangen werden. Im eindimensionalen Fall finden wir dagegen immer ein $x \in U$ enthaltendes offenes Intervall, dessen Rand ganz im Komplement $\mathbb{R} - U$ liegt.

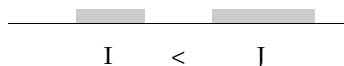
Jede Intervallzerlegung $\mathcal{I}$ besitzt eine natürliche lineare Ordnung:

Definition (*Intervallordnung auf $\mathcal{I}$*)

Sei $\mathcal{I}$ eine Intervallzerlegung. Wir setzen für $I, J \in \mathcal{I}$:

$I < J$ falls $\sup(I) \leq \inf(J)$.

$<$ heißt die *Intervallordnung auf $\mathcal{I}$* .



Die Intervallordnung auf $\mathcal{I}$ ist eine lineare Ordnung auf $\mathcal{I}$.

Der Fall $\sup(I) = \inf(J)$ ist möglich: Sei $U = \mathbb{R} - \{0\}$. Dann ist

$\{]-\infty, 0[,]0, +\infty[\}$

die Intervallzerlegung von U .

Eine Intervallzerlegung haben wir implizit bereits kennengelernt: Ist U das Komplement der Cantormenge im Einheitsintervall, also $U = [0, 1] - C$, so besteht die Intervallzerlegung von U gerade aus den mittleren Drittelintervallen, die bei der Konstruktion der natürlichen Approximation $\langle C_n \mid n \in \mathbb{N} \rangle$ der Cantormenge entfernt wurden.

Nach diesen Vorüberlegungen können wir nun zeigen, daß alle beschränkten, abgeschlossenen und überabzählbaren Teilmengen von $\mathbb{R}$, die keine Intervalle enthalten, nach Entfernung ihrer verborgen isolierten Punkte den Ordnungstyp der Cantormenge besitzen. Dies rechtfertigt dann im nachhinein die folgende Definition:

Definition (*C-artige Mengen*)

Ein $P \subseteq \mathbb{R}$ heißt *C-artig*, falls gilt:

- (i) P ist perfekt und nichtleer,
- (ii) P ist beschränkt,
- (iii) P enthält kein Intervall $[a, b]$ mit $a, b \in \mathbb{R}, a < b$.

Wir zeigen nun, daß die von einer C-artigen Menge gebildeten Zwischenräume eine Ordnung des Typs η der rationalen Zahlen bilden.

Satz

Sei $P \subseteq \mathbb{R}$ eine C-artige Menge, und sei $\mathcal{I}$ die Intervallzerlegung von $U = [\inf(P), \sup(P)] - P$. Weiter seien L und R die Menge der linken bzw. rechten Grenzpunkte der Intervalle aus $\mathcal{I}$, und es sei $E = L \cup R$. Dann gilt:

- (i) o. t. $\langle \mathcal{I}, < \rangle = \eta$.
- (ii) $E' = P$.
- (iii) Sei $x \in P$, $x \neq \inf(P)$. Dann gilt: $x = \sup(\{a \in E \mid a < x\})$ gdw $x \notin R$.
- (iv) Sei $x \in P$, $x \neq \sup(P)$. Dann gilt: $x = \inf(\{a \in E \mid x < a\})$ gdw $x \notin L$.

Beweis

zu (i): Wir verwenden den Charakterisierungssatz von Cantor für den Ordnungstyp η der rationalen Zahlen. Zunächst ist $\mathcal{I}$ abzählbar. Weiter hat $\langle \mathcal{I}, < \rangle$ wegen P perfekt keine Endpunkte und ist also unbeschränkt. $\langle \mathcal{I}, < \rangle$ ist dicht, da andernfalls P ein nichttriviales Intervall oder einen isolierten Punkt enthalten würde. Also gilt o. t. $\langle \mathcal{I}, < \rangle = \eta$.

zu (ii): Seien $a, b \in \mathbb{R}$ und $\inf(P) \leq a < b \leq \sup(P)$.

Dann gilt:

$$(+)\]a, b[\cap E = \emptyset \text{ folgt }]a, b[\cap P = \emptyset.$$

Beweis von (+)

Da P keine Intervalle enthält, existiert ein $z \in]a, b[$ mit $z \in U$.

Sei dann $I \in \mathcal{I}$ mit $z \in I$.

Wegen $]a, b[\cap E = \emptyset$ gilt dann $]a, b[\subseteq I$.

Wegen $I \subseteq U$ also $]a, b[\cap P = \emptyset$. Dies zeigt (+).

Aus (+) folgt, daß jedes $x \in P$ Häufungspunkt von E ist, denn *andernfalls* existieren $y, z \in \mathbb{R}$ mit $y < x < z$ und $]y, x[\cap P =]x, z[\cap P = \emptyset$, und dann ist x ein isolierter Punkt von P , *Widerspruch!*

zu (iii): Die Richtung von links nach rechts ist trivial.

Sei also $x \in P$, $x \neq \inf(P)$, und sei $x \notin R$.

Annahme $\sup(\{a \in E \mid a < x\}) < x$.

Dann existiert nach (+) ein $y < x$ mit $\inf(P) < y$ und $]y, x[\cap P = \emptyset$.

Dann ist aber $]y, x[\subseteq U$, und wegen $x \in P$ ist dann notwendig $x \in R$, *Widerspruch*.

– zu (iv): Analog zu (iii).

Aus diesem Resultat über den Ordnungstyp der Zwischenräume von C-artigen Mengen folgt nun, daß C-artige Mengen selbst immer den Ordnungstyp ξ der Cantormenge haben.

Satz (*Ordnungstyp der C-artigen Mengen*)

Sei $P \subseteq \mathbb{R}$ eine C-artige Menge.

Dann ist o. t. $\langle P, < \rangle = \xi$.

Beweis

Sei $\mathcal{I}$ die Intervallzerlegung von $[\inf(P), \sup(P)] - P$, und sei $\mathcal{J}$ die Intervallzerlegung von $[0, 1] - C$. Weiter seien E und F die Mengen der Grenzpunkte der Intervalle in $\mathcal{I}$ bzw. $\mathcal{J}$.

Nach dem Satz oben existiert ein Ordnungsisomorphismus $f: \mathcal{I} \rightarrow \mathcal{J}$ (bzgl. der Intervallordnungen auf $\mathcal{I}$ und $\mathcal{J}$).

Wir definieren einen Ordnungsisomorphismus $g: E \rightarrow F$ durch

$$g(a) = a', g(b) = b' \text{ falls }]a, b[\in \mathcal{I} \text{ und } f(]a, b[) =]a', b'[.$$

Weiter definieren wir $h: P \rightarrow C$ durch

$$h(x) = \begin{cases} g(x), & \text{falls } x \in E, \\ \inf(C), & \text{falls } x = \inf(P), \\ \sup(\{g(a) \mid a \in E, a < x\}), & \text{falls } x \in P - E, x \neq \inf(P). \end{cases}$$

Seien $x, y \in P, x < y$. Existieren $z_1, z_2 \in E$ mit $x < z_1 < z_2 < y$, so gilt

$$h(x) \leq h(z_1) = g(z_1) < g(z_2) = h(z_2) \leq h(y).$$

Andernfalls gilt aber $]x, y[\in \mathcal{I}$, also $x, y \in E$, und dann ist

$$h(x) = g(x) < g(y) = h(y).$$

Also ist h ordnungstreu und injektiv.

h ist aber auch surjektiv:

Sei hierzu $y \in C$. Ist $y \in F$ oder $y = \inf(C)$, so ist offenbar $y \in \text{rng}(h)$.

Andernfalls gilt nach dem Satz oben $y = \sup(\{a \in F \mid a < y\})$.

Sei dann

$$x = \sup(\{a \in E \mid g(a) < y\}).$$

Dann ist $h(x) = y$.

Also ist $h: P \rightarrow C$ ein Ordnungsisomorphismus, und damit gilt

$$\text{— o. t. } \langle P, < \rangle = \text{o. t. } \langle C, < \rangle = \xi.$$

Hieraus erhalten wir nun leicht:

Korollar (*Mächtigkeit der perfekten Mengen*)

Sei $P \subseteq \mathbb{R}$ perfekt und nichtleer.

Dann gilt $|P| = |\mathbb{R}|$.

Beweis

Enthält P ein Intervall $I = [a, b], a < b$, so ist $|P| = |\mathbb{R}|$ wegen $|I| = |\mathbb{R}|$.

Andernfalls seien $x \in P$ und $a, b \in \mathbb{R}$ mit $a, b \notin P$ und $a < x < b$.

Dann ist $P \cap [a, b]$ C-artig, also ordnungsisomorph zur Cantormenge C .

$$\text{— Insbesondere also } |P \cap [a, b]| = |C| = |\mathbb{R}|, \text{ und damit } |P| = |\mathbb{R}|.$$

Korollar (*Mächtigkeiten der abgeschlossenen Mengen*)

Sei $P \subseteq \mathbb{R}$ abgeschlossen.

Dann ist P abzählbar oder es gilt $|P| = |\mathbb{R}|$.

Beweis

Sei P abgeschlossen und überabzählbar.

Sei $P = P^* \cup A$ die Cantor-Bendixson-Zerlegung von P .

- Dann ist $P^* \subseteq P$ nichtleer und perfekt. Also $|P^*| = |\mathbb{R}|$, und damit $|P| = |\mathbb{R}|$.

Cantor (1884b): „Wir haben also den folgenden Satz:

Eine unendliche abgeschlossene lineare Punktmenge hat entweder die erste Mächtigkeit [ist abzählbar unendlich] oder sie hat die Mächtigkeit des Linearkontinuums ...

Daß dieser merkwürdige Satz eine weitere Gültigkeit auch für *nicht abgeschlossene* lineare Punktmenge und ebenso auch für alle n -dimensionalen Punktmenge hat, wird in späteren Paragraphen bewiesen werden ...

Hieraus wird ... geschlossen werden, daß das *Linearkontinuum die Mächtigkeit der zweiten Zahlenklasse (II.) hat* [d. h. $|\mathbb{R}| = \aleph_1$].

Halle, 15. Novbr. 1883.

(Fortsetzung folgt).“

Cantor hat hier also die Lösung des Kontinuumsproblems angekündigt. Und aus dem Satz über die Mächtigkeit der abgeschlossenen Teilmengen von $\mathbb{R}$ kann man in der Tat die Hoffnung schöpfen, das Kontinuumsproblem lösen zu können. Cantor hatte folgendes Ziel im Auge:

Konstruiere eine beliebige abgeschlossene Menge P der kleinsten überabzählbaren Mächtigkeit $\aleph_1$!

Dann hat P als überabzählbare abgeschlossene Menge einen nichtleeren perfekten Kern, also gilt $|P| = |\mathbb{R}|$. Andererseits gilt nach Konstruktion $|P| = \aleph_1$. Also ist die Mächtigkeit des Kontinuums kleinstmöglich, und die Kontinuums-hypothese ist bewiesen.

Die Konstruktion einer abgeschlossenen Menge $P \subseteq \mathbb{R}$ der Mächtigkeit $\aleph_1$ erscheint nicht unmöglich. Man wirft überabzählbar viele Punkte auf die Linie $\mathbb{R}$, sodaß möglichst wenig Häufungspunkte entstehen, die man wegen der Abgeschlossenheitsforderung immer mit dazu nehmen muß. Um isolierte Punkte braucht man sich nicht zu kümmern ...

Es wäre ein schöner, trickreicher, wahrhaft Cantorscher Beweis der Kontinuums-hypothese gewesen, und ein krönender Höhepunkt eines Lebenswerks. Es ist eine Schande, daß es nicht funktioniert!

Die Mächtigkeiten abgeschlossener Mengen: zweiter Beweis

Obige Darstellung entspricht im wesentlichen Cantors originalem Beweis von 1884, daß überabzählbare abgeschlossene Teilmengen von $\mathbb{R}$ gleichmächtig zu $\mathbb{R}$ sind. Wir geben noch einen weiteren Beweis dieses Satzes, der die Ergebnisse über den Ordnungstyp $\mathbb{R}$ heranzieht und den Begriff der perfekten Mengen nicht verwendet.

Hierzu zunächst zwei für sich interessante Hilfssätze. Der erste zeigt, daß jede überabzählbare Teilmenge A von $\mathbb{R}$ einen Punkt besitzt, der A in einen überabzählbaren linken und einen überabzählbaren rechten Teil zerlegt. Der zweite Satz nutzt dann dieses Resultat, um eine Teilmenge Q von A zu konstruieren, die dicht ist unter der von den reellen Zahlen ererbten Ordnung.

Satz (Aufspaltung überabzählbarer Teilmengen von $\mathbb{R}$)

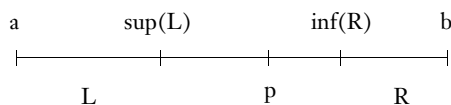
Seien $a, b \in \mathbb{R}$, und sei $P \subseteq [a, b]$ überabzählbar. Dann existiert ein $p \in P$ derart, daß $[a, p] \cap P$ und $[p, b] \cap P$ überabzählbar sind.

Beweis

Wir setzen:

$$L = \{x \in \mathbb{R} \mid a \leq x, [a, x] \cap P \text{ ist abzählbar}\},$$

$$R = \{x \in \mathbb{R} \mid x \leq b, [x, b] \cap P \text{ ist abzählbar}\}.$$



Ist nun $\langle x_n \mid n \in \mathbb{N} \rangle$ eine streng monoton wachsende Folge mit $x_n \in L$ für alle $n \in \mathbb{N}$, und ist $x = \sup(\{x_n \mid n \in \mathbb{N}\})$, so ist $[a, x[\cap P = \bigcup_{n \in \mathbb{N}} [a, x_n] \cap P$ abzählbar, und damit $x \in L$. Also gilt $\sup(L) \in L$. Analog ist $\inf(R) \in R$.

Dann existiert aber ein $p \in P$ mit $\sup(L) < p < \inf(R)$, denn andernfalls wäre $P = ([a, \sup(L)] \cap P) \cup ([\inf(R), b] \cap P)$ abzählbar.

- Wegen $p \notin L, R$ sind dann $[a, p] \cap P$ und $[p, b] \cap P$ überabzählbar.

Satz (Existenz abzählbarer dichter Teilordnungen)

Sei $P \subseteq \mathbb{R}$ überabzählbar.

Dann existiert ein $Q \subseteq P$ mit o.t. $\langle\langle Q, < \rangle\rangle = \eta$.

Beweis

Seien $a, b \in \mathbb{R}$ derart, daß $[a, b] \cap P$ überabzählbar ist (a, b existieren, da sonst $P = \bigcup_{n \in \mathbb{N}} P \cap [-n, n]$ abzählbar wäre).

Wir definieren rekursiv endliche Mengen $Q_n \subseteq P \cup \{a, b\}$ für $n \in \mathbb{N}$, die insbesondere die Eigenschaft haben:

(+) Sind $x, y \in Q_n$, $x < y$, so ist $[x, y] \cap P$ überabzählbar.

Sei $Q_0 = \{a, b\}$. Dann gilt (+) für Q_0 .

Sei Q_n definiert für ein $n \in \mathbb{N}$, und sei $Q_n = \{x_0, \dots, x_k\}$ mit $x_0 < \dots < x_k$.

Nach (+) und dem Satz oben

können wir für jedes $i \in \mathbb{N}$ mit $0 \leq i < k$ ein $p_i \in P$ wählen derart, daß $[x_i, p_i] \cap P$ und $[p_i, x_{i+1}] \cap P$ überabzählbar sind.

Wir setzen dann

$$Q_{n+1} = Q_n \cup \{p_0, \dots, p_{k-1}\}.$$

Dann gilt (+) für Q_{n+1} .

Sei nun $Q = \bigcup_{n \in \mathbb{N}} Q_n - \{a, b\}$.

- Dann ist $Q \subseteq P$ abzählbar, unbeschränkt und dicht, also o. t. $(\langle Q, < \rangle) = \eta$.

Da die Dedekind-Vervollständigung einer Ordnung des Typs η die Mächtigkeit der reellen Zahlen besitzt, erhalten wir hieraus:

Korollar (Mächtigkeiten der abgeschlossenen Mengen)

Sei $P \subseteq \mathbb{R}$ abgeschlossen und überabzählbar.

Dann existiert ein $R \subseteq P$ mit o. t. $(\langle R, < \rangle) = \theta$.

Insbesondere gilt also $|P| = |\mathbb{R}|$.

Beweis

Wegen P überabzählbar existiert ein $Q \subseteq P$ mit o. t. $(\langle Q, < \rangle) = \eta$.

Sei $\langle \mathcal{D}(Q), < \rangle$ die Dedekind-Vervollständigung von $\langle Q, < \rangle$.

Dann gilt o. t. $(\langle \mathcal{D}(Q), < \rangle) = \theta$ (vgl. Kap. 10), also $|\mathcal{D}(Q)| = |\mathbb{R}|$.

Wir identifizieren einen Schnitt (L, R) in $\langle Q, < \rangle$ mit $\sup(L) \in \mathbb{R}$.

Dann ist $\mathcal{D}(Q) \subseteq \mathbb{R}$. Wegen P abgeschlossen und $Q \subseteq P$ ist dann aber

- $\mathcal{D}(Q) \subseteq P$, und dann ist $R = \mathcal{D}(Q)$ wie gewünscht.

Damit haben wir auf zwei verschiedenen Wegen die möglichen Mächtigkeiten der abgeschlossenen Mengen bestimmt. Der erste Weg analysierte die Cantormenge im Detail und ergab ein klares Bild über die Natur nirgends dichter perfekter Mengen. Will man lediglich $|P| = |\mathbb{R}|$ für nichtleere perfekte P zeigen, so kann man den ersten Weg wie folgt kondensieren:

Übung

Sei P nichtleer und perfekt. Zeigen Sie $|P| = |\mathbb{R}|$, indem Sie die Cantormenge ordnungstreu in P einbetten.

[Konstruieren Sie rekursiv abgeschlossene nichtleere Intervalle I_s für endliche 0-2-Folgen $s: \bar{n} \rightarrow \{0, 2\}$ mit:

$$I_{s \smallfrown 0}, I_{s \smallfrown 2} \subseteq I_s, \quad \sup(I_{s \smallfrown 0}) < \inf(I_{s \smallfrown 2}), \quad I_s \cap P \neq \emptyset, \quad \sup(I_s) - \inf(I_s) \leq 1/2^{|\text{dom}(s)|},$$

wobei $s \smallfrown i = s \cup \{(n, i)\}$ für $i \in \{0, 2\}$ die Verlängerung von $s: \bar{n} \rightarrow \{0, 2\}$ um i ist.

Betrachte dann $x_g = \bigcap_{n < \omega} I_{g \upharpoonright \bar{n}} \in P$ für $g: \mathbb{N} \rightarrow \{0, 2\}$.]

Eine Verfeinerung der Konstruktion liefert wieder einen Ordnungsisomorphismus zwischen der Cantormenge und einer nirgends dichten nichtleeren perfekten Menge P .

Perfekt reduzierbare Mengen

Eine weitere Strategie zur Lösung des Kontinuumsproblems betrifft die Frage nach der Existenz großer abgeschlossener Teilmengen in großen Teilmengen der reellen Zahlen.

Für alle Teilmengen P von $\mathbb{R}$ ist $P - P'$ abzählbar. Überabzählbare Teilmengen von $\mathbb{R}$ enthalten also sehr viele Häufungspunkte. Enthalten überabzählbare Mengen vielleicht sogar stets eine überabzählbare abgeschlossene Menge? Wenn ja, so enthalten sie sogar eine überabzählbare perfekte Teilmenge. Diese Frage ist 1884 von Ludwig Scheeffter (1859 – 1885) gestellt worden.

Definition (*perfekt reduzierbar, Scheeffter-Eigenschaft*)

- (i) $A \subseteq \mathbb{R}$ heißt *perfekt reduzierbar*, falls ein nichtleeres perfektes $P \subseteq A$ existiert.
- (ii) $A \subseteq \mathbb{R}$ hat die *Scheeffter-Eigenschaft*, falls A abzählbar oder perfekt reduzierbar ist.

*Scheeffter (1884): „Der Nachweis der Existenz [gewisser Funktionen] ist uns indes nicht in voller Allgemeinheit, sondern nur unter der Voraussetzung gelungen, daß die Menge P entweder selbst *perfekt* ist oder doch einen *perfekten* Bestandteil enthält. Die von uns aufgeworfene Frage [ein Problem aus der Differentialrechnung einer reellen Veränderlichen] wird also im Folgenden nicht vollständig erledigt; denn es bleibt der Fall übrig, daß die Menge P von höherer als der ersten Mächtigkeit [nicht abzählbar] ist, ohne einen perfekten Bestandteil zu enthalten. Freilich ist es zweifelhaft, ob solche Mengen überhaupt vorkommen können; Beispiele dieser Art dürften wenigstens bisher nicht bekannt sein. Sollte es sich daher in der weiteren Entwicklung der Mengenlehre herausstellen, daß wirklich jede Punktmenge, deren Mächtigkeit höher als diejenige der abzählbaren Mengen ist, einen perfekten Bestandteil enthält, so würde damit gleichzeitig der letzte Schritt zur Beantwortung unserer Frage getan sein. Bei dem gegenwärtigen Standpunkte der Theorie muß jedenfalls die Bedingung, daß die nicht abzählbare Menge P einen perfekten Bestandteil enthalten solle, ausdrücklich angegeben werden.“*

Scheeffter verstarb 1885 an Typhus. Der in Königsberg geborene und zuletzt in München tätige Mathematiker hatte sehr früh die Mengenlehre Cantors in der Analysis angewendet (und bildet zusammen mit Poincaré und Mittag-Leffler hierin ein Pioniertrio). Der Tod des 26-jährigen markiert einen tragischen Verlust in der Geschichte der Mengenlehre. Cantor hat noch im gleichen Jahr einen Nekrolog auf den Menschen und den Mathematiker verfaßt [Cantor 1932].

Nach dem Satz von Cantor-Bendixson haben alle abgeschlossenen Mengen die Scheeffter-Eigenschaft. Für nichtleere offene Mengen gilt die perfekte Reduzierbarkeit trivialerweise, denn diese Mengen enthalten ein nichttriviales abgeschlossenes Intervall.

Mit dem Nachweis, daß jede Teilmenge von $\mathbb{R}$ die Scheeffter-Eigenschaft hat, hätten wir die Kontinuumshypothese bewiesen: Denn sei $A \subseteq \mathbb{R}$ überabzählbar.

Da A perfekt reduzierbar ist, enthält A eine nichtleere perfekte Teilmenge P . Wegen $|P| = |\mathbb{R}|$ ist also $|A| = |\mathbb{R}|$.

Wir werden zeigen, daß es überabzählbare Teilmengen von $\mathbb{R}$ gibt, die nicht perfekt reduzierbar sind. Eine solche Menge wird aber sehr abstrakt durch Diagonalisierung konstruiert – und sie *muß* sehr abstrakt konstruiert werden, denn es läßt sich (innerhalb der deskriptiven Mengenlehre) zeigen, daß alle „vernünftigen“, natürlich auftretenden Teilmengen von $\mathbb{R}$ tatsächlich die Scheeffer-Eigenschaft besitzen, weshalb Scheeffer auch keine „Beispiele“ für solche Ausnahmемengen in der Analysis fand.

Die Idee der diagonalen Konstruktion einer überabzählbaren, nicht perfekt reduzibaren Menge ist: Mit Hilfe einer Aufzählung aller perfekten Mengen konstruieren wir eine überabzählbare Menge, die von jeder Menge der Aufzählung mindestens einen Punkt ausläßt, und damit keine Menge der Aufzählung als Teilmenge enthalten kann.

Um die Konstruktion durchführen zu können, müssen wir die Mächtigkeit der Menge aller perfekten Mengen bestimmen. Der Weg hierzu führt über die offenen Mengen.

Satz (die Mächtigkeit der Menge aller offenen Mengen)

Sei $\mathcal{G} = \{U \subseteq \mathbb{R} \mid U \text{ ist offen}\}$. Dann gilt $|\mathcal{G}| = |\mathbb{R}|$.

Beweis

Sei $I_0, I_1, I_2, \dots$ eine Aufzählung der rationalen Intervalle.

Wir definieren eine Funktion $f: \mathcal{G} \rightarrow \mathcal{P}(\mathbb{N})$ durch:

$$f(U) = \{n \in \mathbb{N} \mid I_n \subseteq U\} \quad \text{für } U \in \mathcal{G}.$$

Dann ist f injektiv, denn für alle $U \in \mathcal{G}$ gilt $U = \bigcup_{n \in f(U)} I_n$.

Also gilt $|\mathcal{G}| \leq |\mathcal{P}(\mathbb{N})| = |\mathbb{R}|$.

Seien nun $J_0, J_1, \dots$ paarweise disjunkte nichtleere offene Intervalle, z. B. $J_n =]n, n+1[$ für $n \in \mathbb{N}$. Wir definieren $g: \mathcal{P}(\mathbb{N}) \rightarrow \mathcal{G}$ durch:

$$g(A) = \bigcup_{n \in A} J_n \quad \text{für } A \subseteq \mathbb{N}.$$

Dann ist g injektiv, also $|\mathcal{P}(\mathbb{N})| \leq |\mathcal{G}|$.

– Insgesamt also $|\mathcal{G}| = |\mathcal{P}(\mathbb{N})|$.

Korollar (die Mächtigkeit der Mengen aller abgeschlossenen und perfekten Mengen)

Seien $\mathcal{F} = \{P \subseteq \mathbb{R} \mid P \text{ ist abgeschlossen}\}$, $\mathcal{P} = \{P \subseteq \mathbb{R} \mid P \text{ ist perfekt}\}$.

Dann gilt $|\mathcal{F}| = |\mathcal{P}| = |\mathbb{R}|$.

Beweis

$f: \mathcal{G} \rightarrow \mathcal{F}$ mit $f(U) = \mathbb{R} - U$ für $U \in \mathcal{G}$ ist bijektiv, also $|\mathcal{F}| = |\mathcal{G}| = |\mathbb{R}|$.

Offenbar gilt $|\mathcal{P}| \leq |\mathcal{F}| = |\mathbb{R}|$.

Sei $J_n =]n, n+1[$ für $n \in \mathbb{N}$. Wir definieren $h: \mathcal{P}(\mathbb{N}) \rightarrow \mathcal{P}$ durch:

$$h(A) = \bigcup_{n \in A} J_n \quad \text{für } A \subseteq \mathbb{N}.$$

Dann ist h injektiv, also $|\mathcal{P}(\mathbb{N})| \leq |\mathcal{P}|$.

– Insgesamt folgt $|\mathcal{P}| = |\mathbb{R}|$.

Obwohl es also $|\mathcal{P}(\mathbb{R})| > |\mathbb{R}|$ viele Teilmengen von $\mathbb{R}$ gibt, existieren jeweils lediglich $|\mathbb{R}|$ viele offene, abgeschlossene und perfekte Mengen. Allgemeiner gilt $|\mathfrak{B}| = |\mathbb{R}|$ für die Borelmengen $\mathfrak{B}$ von $\mathbb{R}$, wie man durch eine Induktion der Länge ω_1 über die Borel-Hierarchie leicht zeigt.

Übung

Ist $|W(\omega_1)| < |\mathbb{R}|$, so existiert ein $X \subseteq \mathbb{R}$ ohne die Scheeffer-Eigenschaft.

Mit Hilfe des Wohlordnungssatzes können wir nun zeigen, daß es eine überabzählbare Teilmenge der reellen Zahlen gibt, die keine perfekte Teilmenge besitzt. Der folgende Satz und sein „diagonaler“ Beweis stammen von Felix Bernstein aus dem Jahr 1908, der damit die ein Vierteljahrhundert zuvor gestellte Frage von Ludwig Scheeffer negativ beantwortet hat.

Satz (Existenz nicht perfekt reduzierbarer Mengen der Größe von $\mathbb{R}$)

Es existiert ein $X \subseteq \mathbb{R}$ mit $|X| = |\mathbb{R}|$, das nicht perfekt reduzierbar ist.

Genauer gilt: Es gibt ein $X \subseteq \mathbb{R}$ mit $|X| = |\mathbb{R} - X| = |\mathbb{R}|$, sodaß sowohl X als auch $\mathbb{R} - X$ nicht perfekt reduzierbar sind.

Beweis

Sei $\kappa = |\mathbb{R}|$ die Kardinalität von $\mathbb{R}$. Nach dem Satz oben über die Mächtigkeit von $\mathcal{P}$ gibt es eine Aufzählung $\langle P_\alpha \mid \alpha < \kappa \rangle$ der nichtleeren perfekten Mengen. Weiter sei $\langle \mathbb{R}, < \rangle$ eine Wohlordnung der reellen Zahlen.

Wir definieren nun rekursiv $x_\alpha, y_\alpha \in P_\alpha$ für $\alpha < \kappa$ wie folgt.

Rekursionsschritt $\alpha < \kappa$

Seien x_β, y_β definiert für $\beta < \alpha$. Wir setzen:

$x_\alpha =$ „das $<$ -kleinste x mit $x \in P_\alpha^1 = P_\alpha - \{x_\beta, y_\beta \mid \beta < \alpha\}$ “,

$y_\alpha =$ „das $<$ -kleinste x mit $x \in P_\alpha^2 = P_\alpha^1 - \{x_\alpha\}$ “.

[Es gilt $P_\alpha^1, P_\alpha^2 \neq \emptyset$ für alle $\alpha < \kappa = |\mathbb{R}|$, denn $|P_\alpha| = |\mathbb{R}|$ für alle $\alpha < \kappa$.]

Wir setzen nun $X = \{x_\alpha \mid \alpha < \kappa\}$, $Y = \{y_\alpha \mid \alpha < \kappa\}$.

Nach Konstruktion gilt dann:

- (i) $X \cap Y = \emptyset$, $|X| = |Y| = |\mathbb{R}| = \kappa$,
- (ii) $X \cap P_\alpha \neq \emptyset$ für alle $\alpha < \kappa$,
- (iii) $Y \cap P_\alpha \neq \emptyset$ für alle $\alpha < \kappa$.

Somit $\text{non}(P_\alpha \subseteq X)$ für alle $\alpha < \kappa$. Also enthält X keine perfekte Menge, und es gilt $|X| = |\mathbb{R}|$. Zudem enthält auch $\mathbb{R} - X \supseteq Y$ keine perfekte Menge, — und wegen $|Y| = |\mathbb{R}|$ gilt $|\mathbb{R} - X| = |\mathbb{R}|$.

Die Konstruktion des Beweises benutzt keine speziellen Eigenschaften von perfekten Mengen. Sie wird durch ein Diagonalisierungsverfahren getragen, das man allgemein durchführen kann, wenn ein Teilmengensystem einer Menge nur aus Mengen der Mächtigkeit von M besteht, und darüber hinaus selbst die Mächtigkeit von M besitzt:

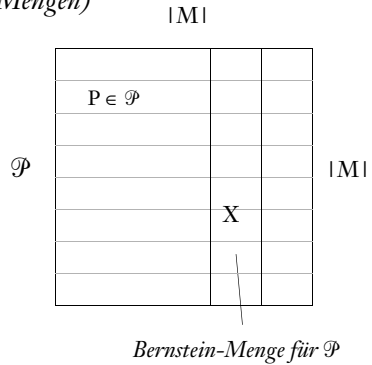
Definition (*uniforme Systeme und Bernstein-Mengen*)

Sei M eine unendliche Menge.
 Eine Menge $\mathcal{P} \subseteq \mathcal{P}(M)$ heißt
uniformes System auf M , falls gilt:

- (i) $|\mathcal{P}| = |M|$,
- (ii) $|P| = |M|$ für alle $P \in \mathcal{P}$.

Ein $X \subseteq M$ heißt eine *Bernstein-Menge* für $\mathcal{P}$, falls gilt:

- (i) $|X| = |M - X| = |M|$,
- (ii) für kein $P \in \mathcal{P}$ gilt $P \subseteq X$
 oder $P \subseteq M - X$.



Die in obigem Beweis konstruierte nicht perfekt reduzierbare Menge ist also eine Bernstein-Menge für das uniforme System der nichtleeren perfekten Teilmengen von $\mathbb{R}$. Allgemein gilt nun:

Satz (*Existenz von Bernstein-Mengen*)

Sei M eine unendliche Menge, und sei $\mathcal{P}$ ein uniformes System auf M .
 Dann existiert eine Bernstein-Menge für $\mathcal{P}$.

Beweis

Analog zum Beweis der Existenz nicht perfekt reduzierbarer Teilmengen
 – von $\mathbb{R}$; wir benutzen eine Wohlordnung von M .

Für ein uniformes System auf M kann man also immer eine Zerlegung von M in zwei Teile finden, sodaß jedes Element des Systems sowohl den einen als auch den anderen Teil trifft.

Teilmengen von $\mathbb{R}$ mit der Scheeffer-Eigenschaft

Wir haben gezeigt, daß alle offenen und alle abgeschlossenen Mengen die Scheeffer-Eigenschaft besitzen. Weiter haben wir eine Menge ohne diese Eigenschaft konstruiert – in sehr abstrakter Weise. Es stellt sich die Frage, welchen Teilmengen von $\mathbb{R}$ die Scheeffer-Eigenschaft noch zukommt. Alle diese Mengen sind dann abzählbar oder von der Mächtigkeit der reellen Zahlen, d.h. (CH) gilt für alle Mengen mit der Scheeffer-Eigenschaft.

In der Tat bestätigt sich die Intuition: Einfache Teilmengen von $\mathbb{R}$ haben die Scheeffer-Eigenschaft. Cantors Resultat über die abgeschlossenen Mengen ist in der Folgezeit mehrfach verallgemeinert worden. Es gilt:

(+) Alle Borelmengen haben die Scheeffer-Eigenschaft.

Diesen Satz haben unabhängig voneinander Hausdorff und Pavel Alexandrov (1896 – 1982) im Jahre 1916 gezeigt. 1903 hatte William Young (1863 – 1942) die Scheeffer-Eigenschaft für $\mathcal{G}_\delta$ -Mengen nachgewiesen, und Hausdorff konnte 1914 den Satz von Young auf $\mathcal{G}_{\delta\sigma\delta}$ -Mengen erweitern.

Für die Borelmengen ist also Cantors Kontinuumshypothese beweisbar richtig. Die Frage ist nun, ob es noch weitere „einfache“ Teilmengen von $\mathbb{R}$ gibt, für die das Resultat von Hausdorff-Alexandrov noch einmal erweitert werden könnte. In der Tat gibt es eine natürliche Klasse von Punktmengen hinter den Borelmengen: An die Borel-Hierarchie schließen sich die sog. projektiven Mengen an. Hier gerät die Scheeffer-Eigenschaft zwar ins Reich der Unabhängigkeit der üblichen Mathematik, dafür aber in den Einfluß von großen Kardinalzahlen. Wir geben im folgenden einen kurzen Überblick über die wichtigsten Definitionen und Ergebnisse.

Die einfachste Erweiterung der Borelmengen um komplexere definierbare Mengen bilden die sog. analytischen Mengen (bzw. dual die koanalytischen Mengen). Wir geben drei äquivalente Definitionen dieser Mengen kurz an.

Man definiert zunächst in naheliegender Weise den Begriff „ $A \subseteq \mathbb{R}^n$ ist abgeschlossen“ für $1 \leq n < \omega$.

Eine Punktmenge $A \subseteq \mathbb{R}^n$ in einem n -dimensionalen Raum ist abgeschlossen, wenn sie alle ihre Häufungspunkte enthält. Ein Häufungspunkt von A ist ein Punkt $x \in \mathbb{R}^n$, für den gilt: Jede n -dimensionale Kugel um x enthält einen Punkt der Menge $A - \{x\}$.

Die offenen Mengen in $\mathbb{R}^n$ sind wieder die Komplemente der abgeschlossenen Mengen in $\mathbb{R}^n$. Analog zum Fall $n = 1$ definiert man dann weiter die Borel-Hierarchie, und erhält so den Begriff „ $A \subseteq \mathbb{R}^n$ ist eine Borelmenge“.

Ist $P \subseteq \mathbb{R}^{n+1}$ eine Punktmenge des $(n+1)$ -dimensionalen Kontinuums, so ist die Projektion von P auf den Raum $\mathbb{R}^n$ definiert durch:

$$\text{Pr}(P) = \{ (x_1, \dots, x_n) \in \mathbb{R}^n \mid \text{es gibt ein } y \text{ mit } (x_1, \dots, x_n, y) \in P \}.$$

Definition (*analytisch*)

Sei $A \subseteq \mathbb{R}$. A heißt *analytisch*, *Suslinsch* oder eine Σ_1^1 -Menge, falls gilt:

$$A = \text{Pr}(B) \text{ für eine Borelmenge } B \subseteq \mathbb{R}^2.$$

Statt „ B Borelsch“ genügt bereits „ B ist $\mathcal{G}_\delta$ “, d. h. B ist ein abzählbarer Schnitt von offenen Mengen des $\mathbb{R}^2$. Zieht man die Theorie über dem Bairerraum statt $\mathbb{R}$ auf, so genügt bereits „ B ist abgeschlossen“.

Äquivalent ist die folgende Definition, die den Umweg über höhere Dimensionen vermeidet:

A ist *analytisch*, falls $A = f'' B$ für eine Borelmenge B und ein stetiges $f: \mathbb{R} \rightarrow \mathbb{R}$.

Die analytischen Mengen sind also die Bilder der Borelmengen unter stetigen Funktionen.

Die analytischen Mengen sind nicht abgeschlossen unter Komplementbildung, und man betrachtet daher:

Definition (*koanalytisch*)

Sei $A \subseteq \mathbb{R}$. A heißt *koanalytisch* oder eine Π_1^1 -Menge, falls $\mathbb{R} - A$ analytisch ist.

Suslin hat 1917 gezeigt, daß die Borelmengen auf $\mathbb{R}$ genau diejenigen Mengen sind, die sowohl analytisch als auch koanalytisch sind (Satz von Suslin). Die Borelmengen lassen sich also aus den analytischen Mengen definieren.

Die analytischen Mengen kann man weiter durch eine originelle Operation aus den abgeschlossenen Mengen von $\mathbb{R}$ gewinnen. Sie geht auf Alexandrov und Suslin zurück.

Definition (*die $\mathcal{A}$ -Operation*)

Sei $S = \{s \mid s : \bar{n} \rightarrow \mathbb{N} \text{ für ein } n < \omega\}$ die Menge der endlichen Folgen von natürlichen Zahlen. Für jedes $s \in S$ sei $A_s \subseteq \mathbb{R}$.

Dann setzen wir:

$$\mathcal{A}(\langle A_s \mid s \in S \rangle) = \bigcup_{f \in {}^{\mathbb{N}}\mathbb{N}} \bigcap_{n \in \mathbb{N}} A_{f \upharpoonright \bar{n}}.$$

Das sieht wilder aus als es ist: Man bestückt alle Knoten des sich an jeder Stelle unendlich verzweigenden Baumes S mit einer Teilmenge von $\mathbb{R}$. Für jeden Pfad $f \in {}^{\mathbb{N}}\mathbb{N}$ durch den Baum wird nun der Schnitt der Mengen gebildet, die auf dem Pfad liegen, und alle so erhaltenen Schnitte – $\mathbb{R}$ -viele – werden vereinigt.

Die Mengen, die man durch die $\mathcal{A}$ -Operation erhalten kann, wenn man den Baum mit abgeschlossenen Teilmengen von $\mathbb{R}$ bestückt, sind nun genau die analytischen Mengen! Man erhält nicht mehr, wenn man den Baum mit Borelmengen bestückt anstatt mit abgeschlossenen Mengen. Auch eine Bestückung mit analytischen Mengen liefert wieder analytische Mengen, d. h. die analytischen Mengen sind abgeschlossen unter der $\mathcal{A}$ -Operation.

Die analytischen und koanalytischen Mengen bilden nur die erste Stufe einer Hierarchie der Länge ω , der sog. projektiven Hierarchie von Nikolai Lusin (1883 – 1950) und Wacław Sierpiński (1882 – 1969) aus dem Jahre 1925, die in der deskriptiven Mengenlehre eine zentrale Position einnimmt.

Definition (*projektive Mengen*)

$A \subseteq \mathbb{R}$ heißt *projektiv*, falls A durch iterierte Projektion und Komplementbildung aus einer Borelmenge $B \subseteq \mathbb{R}^n$ für ein $1 \leq n < \omega$ gewonnen werden kann.

Es ergibt sich eine Hierarchie der Länge ω durch Zählen von Projektionen und Komplementbildungen. Für $A \subseteq \mathbb{R}^n$ sei $C(A) = \mathbb{R}^n - A$. Dann setzt man:

$$A \in \Sigma_1^1, \quad \text{falls } A = \text{Pr}(B),$$

$$A \in \Pi_1^1, \quad \text{falls } A = C(\text{Pr}(B)),$$

$$A \in \Sigma_2^1, \quad \text{falls } A = \text{Pr}(C(\text{Pr}(B))),$$

$$A \in \Pi_2^1, \quad \text{falls } A = C(\text{Pr}(C(\text{Pr}(B)))),$$

$$A \in \Sigma_3^1, \quad \text{falls } A = \text{Pr}(C(\text{Pr}(C(\text{Pr}(B))))), \quad \text{usw.,}$$

für eine Borelmenge B eines geeignet dimensionalen Raumes $\mathbb{R}^n$. A ist also eine projektive Menge, falls $A \in \bigcup_{1 \leq n < \omega} \Sigma_n^1 (= \bigcup_{1 \leq n < \omega} \Pi_n^1)$. Die Hierarchie ist echt, es gilt $\Sigma_n^1 \cup \Pi_n^1 \subset \Sigma_{n+1}^1 \cap \Pi_{n+1}^1$ für alle $1 \leq n < \omega$.

Die Definition liefert natürlich allgemeiner einen Begriff von „ $A \subseteq \mathbb{R}^k$ ist projektiv“, und eine zugehörige Hierarchie. So ist etwa $A \subseteq \mathbb{R}^3$ analytisch, falls $A = \text{Pr}(B)$ für eine Borelmenge $B \subseteq \mathbb{R}^4$, usw.

Alternativ kann man die projektive Hierarchie auch durch „Bild unter einer stetigen Funktion“ und Komplementbildung aufbauen. Induktiv definiert man für $n \geq 1$: $\Sigma_{n+1}^1 = \{ f''A \mid A \in \Pi_n^1, f: \mathbb{R} \rightarrow \mathbb{R} \text{ stetig} \}$, $\Pi_{n+1}^1 = \{ \mathbb{R} - A \mid A \in \Sigma_{n+1}^1 \}$.

Zurück zur Scheeffe-Eigenschaft, und damit letztendlich zu den großen Kardinalzahlen! Alexandrov und Suslin haben 1917 gezeigt:

(++) Die analytischen Mengen haben die Scheeffe-Eigenschaft.

Dagegen kann man in der üblichen Mathematik nicht mehr zeigen, daß auch alle koanalytischen Mengen die Scheeffe-Eigenschaft haben! Gödel zeigte nämlich, daß es in seinem Modell L eine überabzählbare koanalytische Menge gibt, die keine nichtleere perfekte Teilmenge besitzt. Dagegen konstruierte Solovay in den 60er Jahren des 20. Jahrhunderts mit Hilfe einer unerreichbaren Kardinalzahl ein Modell, in welchem sogar alle projektiven Mengen die Scheeffe-Eigenschaft besitzen.

Gültigkeit in Modellen ist hübsch, besser sind aber direkte Implikationen. Eines der großen Resultate der jüngeren Mengenlehre ist nun das Martin-Steel-Resultat (1989). Die Aussage kann man sehr grob etwa so formulieren:

(+++) Existieren genügend starke, aber immer noch gut verstandene große Kardinalzahlen, so haben die projektiven Mengen alle erdenklichen Regularitätseigenschaften.
Insbesondere hat jede projektive Menge die Scheeffe-Eigenschaft, und ist also abzählbar oder gleichmächtig mit $\mathbb{R}$.

Der Leser vergleiche dies mit Gödels Resultat über L , wo bereits für eine koanalytische Menge die Scheeffe-Eigenschaft verletzt ist. Mit der Trommel hervorzuheben ist, daß der Satz nicht über Modelle redet: Es handelt sich um eine direkte Implikation, nicht um eine relative Gültigkeit in einem Modell, und darin unterscheidet sich (+++) wesentlich etwa von dem Ergebnis von Solovay, das nur die relative Widerspruchsfreiheit der Hypothese „alle projektiven Mengen haben die Scheeffe-Eigenschaft“ liefert.

„Genügend stark“ ist von Donald Martin und John Steel genau bestimmt worden, die Annahme lautet in etwa „es existieren unendlich viele Woodin-Kardinalzahlen“. Eine meßbare Kardinalzahl liefert bereits Regularität (insb. die Scheeffe-Eigenschaft) für Σ_2^1 -Mengen (Martin 1970). Für höhere Komplexitäten muß man dann die viel größeren Woodin-Kardinalzahlen bemühen. Es gibt eine übergeordnete Super-Regularitätseigenschaft, die sog. Determiniertheit von Teilmengen von $\mathbb{R}$. „Projektive Determiniertheit“ besagt, daß alle projektiven Mengen diese Eigenschaft haben, und projektive Determiniertheit folgt aus der Existenz von unendlich vielen Woodin-Kardinalzahlen.

Neben „ A hat die Scheeffe-Eigenschaft“ gibt es noch zwei weitere prominente Regularitätseigenschaften von Teilmengen von $\mathbb{R}$: Die eine ist „ A ist Lebesgue-meßbar“, d.h. „es gibt eine Borelmenge B mit $A \Delta B = (A - B) \cup (B - A)$ hat Längenmaß 0“ (vgl. die Defini-

tion oben), und „A hat die Baire-Eigenschaft“, d.h. „es gibt ein offenes U mit $A \Delta U$ ist mager“. Es zeigt sich: $A \subseteq \mathbb{R}$ ist Lebesgue-meßbar *gdw* $A = F \cup N$ für eine $\mathcal{F}_\sigma$ -Menge F und eine Nullmenge N ; $A \subseteq \mathbb{R}$ hat die Baire-Eigenschaft *gdw* $A = G \cup M$ für eine $\mathcal{G}_\delta$ -Menge G und eine magere Menge M .

Die Lebesgue-meßbaren Mengen und die Mengen mit der Baire-Eigenschaft bilden jeweils eine σ -Algebra. Diese σ -Algebren enthalten die analytischen und koanalytischen Mengen (Lusin 1917). Lebesgue-Meßbarkeit und Baire-Eigenschaft für höhere Stufen der projektiven Hierarchie verhalten sich wie die Scheeffer-Eigenschaft, folgen also insbesondere aus der Existenz genügend großer Kardinalzahlen, und gelten nicht in Gödels Modell L . Die Gegenbeispiele in L mit kleinstmöglicher Komplexität sind dabei Δ_2^1 -Mengen, wobei $\Delta_2^1 = \Pi_2^1 \cap \Sigma_2^1$.

Eine Bernstein-Menge X für $\mathcal{P} = \{P \subseteq \mathbb{R} \mid P \text{ nichtleer und perfekt}\}$ ist weder meßbar noch hat sie die Baire-Eigenschaft: Die $\mathcal{F}_\sigma$ - und $\mathcal{G}_\delta$ -Mengen haben die Scheeffer-Eigenschaft, und abzüglich einer Nullmenge bzw. einer mageren Menge sind ein meßbares bzw. ein Bairesches $X \subseteq \mathbb{R}$ und sein Komplement $\mathcal{F}_\sigma$ oder $\mathcal{G}_\delta$, X oder $\mathbb{R} - X$ enthält also eine nichtleere perfekte Teilmenge, falls X meßbar ist oder die Baire-Eigenschaft hat.

Borel-Determiniertheit (also Super-Regularität für die Borelmengen) läßt sich noch ohne große Kardinalzahlen zeigen (Martin 1975), und liefert dann ein weitgespanntes Dach über den Sätzen von Hausdorff und Alexandrov betreffend die Scheeffer-Eigenschaft der Σ_1^1 -Mengen (und betreffend die Lebesgue-Meßbarkeit und Baire-Eigenschaft der Σ_1^1 - und Π_1^1 -Mengen). Für die Borel-Determiniertheit muß allerdings die mengentheoretische Maschinerie schon ungewöhnlich stark beansprucht werden (Friedman 1971): Es werden die Mengen $\mathcal{P}^\alpha(\mathbb{R})$ für $\alpha < \omega_1$ benötigt, um dieses Resultat über Teilmengen von $\mathbb{R}$ zeigen zu können.

Zusammenfassend halten wir fest:

- (1) In der üblichen Mathematik läßt sich zeigen: Alle Borelmengen sind determiniert. Alle analytischen Mengen haben die Scheeffer-Eigenschaft. Alle analytischen und koanalytischen Mengen sind Lebesgue-meßbar und haben die Baire-Eigenschaft.
- (2) Die Resultate in (1) sind bestmöglich: In Gödels L gibt es eine koanalytische Menge ohne die Scheeffer-Eigenschaft, und eine Δ_2^1 -Menge (also eine Menge, die sowohl Σ_2^1 als auch Π_2^1 ist), die nicht Lebesgue-meßbar ist und nicht die Baire-Eigenschaft besitzt.
- (3) Determiniertheit (und damit alle erdenkliche Regularität) für alle projektiven Mengen folgt aus der Existenz großer Kardinalzahlen (und darüber hinaus auch aus vielen kombinatorischen Prinzipien, etwa aus „ $\mathcal{C}_{\omega_1}$ ist ω_1 -dicht“).

Dies ist alles eine längere Geschichte, die zwar in der Gründerzeit der Mengenlehre beginnt, aber lange Jahre über sie hinaus weitergesponnen wurde. Wir belassen es hier also bei diesem kleinen und notgedrungen unhistorischen Ausblick. Die reellen Zahlen bilden eine komplizierte Struktur, die wir nur in Ansätzen verstehen – dies ist vielleicht die Grundbotschaft der Mengenlehre. Zum ersten Mal wurde dies deutlich durch Cantors Punktmengenanalyse, seine Entdeckung der transfiniten Prozesse und seine Kontinuumsannahme. Das den analytischen Rechenmeistern seit Jahrhunderten so vertraut erscheinende Objekt $\mathbb{R}$ erweist sich in der Mengenlehre als ein Abgrund. Tief ist der Brunnen der reellen Zahlen. Sollte man ihn nicht unerforschbar nennen? Zuweilen erscheint er als ein unerhört waghalsiges Konstrukt.

Das Buch „Reelle Zahlen. Das klassische Kontinuum und die natürlichen Folgen“ ([Deiser 2007]) greift viele dieser Gesichtspunkte auf. Dort findet sich auch eine Einführung in unendliche Spiele und ein Beweis der Borel-Determiniertheit.

Arthur Schoenflies über nicht perfekt reduzierbare Mengen

„L. Scheeffer hat bereits die Frage aufgeworfen, ob jede nicht abzählbare Menge einen perfekten Bestandteil besitzen müsse [Fußnote hierzu: Acta math. 5 (1884), S. 287.] Dies trifft jedoch nicht zu, wie von F. Bernstein nachgewiesen wurde [Fußnote: Leipz. Ber. 60 (1908), S. 325.] Mengen ohne perfekten Bestandteil bezeichnet er als total imperfekte Mengen. Ihre Existenz ist für den Fall, daß $c [= |\mathbb{R}|] > \aleph_1$ ist, evident. Denn da es Teilmengen des Kontinuums von der Mächtigkeit $\aleph_1$ gibt, und jede perfekte Menge die Mächtigkeit c besitzt, so kann eine Teilmenge der Mächtigkeit $\aleph_1$ keinen perfekten Bestandteil enthalten, ist also total imperfekt.

Für den Fall, daß $c = \aleph_1$ ist, folgt der Bernsteinsche Satz aus dem S. 210 mitgeteilten Theorem [entspricht dem Satz über die Existenz von Bernstein-Mengen], das er über wohlgeordnete Mengen bewiesen hat. Wir haben dazu das Kontinuum als wohlgeordnete Menge zugrunde zu legen; es stellt dann die Menge W des eben genannten Theorems dar. Jede perfekte Menge ist eine Teilmenge von W , und da sie ebenfalls die Mächtigkeit c hat, so ist sie eine Menge w_α [wobei $\langle w_\alpha \mid \alpha < c \rangle$ eine Aufzählung der nichtleeren perfekten Teilmengen von $\mathbb{R}$ ist]. Endlich hat aber auch die Menge aller perfekten Punktmengen des Raumes ... die Mächtigkeit c . Damit sind die Bedingungen des genannten Theorems realisiert, und die Menge A [entspricht X oben], die wir S. 210 aufgestellt haben, ist notwendig total imperfekt, da sie keine Menge w_α ganz enthält. Sie hat selbst die Mächtigkeit c .“

(Arthur Schoenflies 1913, „Entwicklung der Mengenlehre und ihrer Anwendungen“, S. 361ff.)

13. Die Vielheit aller Ordinalzahlen

Zum Abschluß dieses Abschnitt zeigen wir, daß die Annahme der Existenz der Menge aller Ordinalzahlen zu Widersprüchen führt. Diese Tatsache ist in der Literatur als Paradoxie von Burali-Forti bekannt [Burali-Forti, 1897]. Cantor war mit der Unmöglichkeit einer Zusammenfassung aller Ordinalzahlen zu einem Ganzen aber spätestens seit 1897 vertraut, wie aus seinen Briefen an Dedekind und Hilbert hervorgeht.

Ordinalzahlparadoxon von Cantor und Burali-Forti

Wir setzen:

$$\Omega = \{ \alpha \mid \alpha \text{ ist eine Ordinalzahl} \} = \\ \{ \alpha \mid \alpha \text{ ist Ordnungstyp einer Wohlordnung} \}.$$

Für die Relation $<$ auf $\Omega \times \Omega$ gilt:

$\langle \Omega, < \rangle$ ist eine Wohlordnung.

Dann hat $\langle \Omega, < \rangle$ einen bestimmten Ordnungstyp. Sei

$$\delta = \text{o.t.}(\langle \Omega, < \rangle).$$

Offenbar gilt

(+) $\alpha < \delta$ für jede Ordinalzahl α .

(Denn $\langle W(\alpha), < \rangle$ ist ein Anfangsstück von $\langle \Omega, < \rangle$ des Ordnungstyps α .)

– Also gilt wegen (+) $\delta < \delta$, *Widerspruch!*

Cantor (3. 8. 1899, Brief an Dedekind): „Es kann Ω' [hier = Ω] (und daher auch Ω) [hier = $\Omega - \{0\}$] keine konsistente Vielheit [Menge] sein; wäre Ω' konsistent, so würde ihr als einer wohlgeordneten Menge eine Zahl δ zukommen, die größer wäre als alle Zahlen des Systems Ω ; im System Ω kommt aber, weil es *alle* Zahlen umfaßt, auch die Zahl δ vor; es wäre also δ größer als δ , was ein Widerspruch ist. Also

A. Das System Ω aller Zahlen ist eine inkonsistente, eine absolut unendliche Vielheit [echte Klasse].“

Die Zusammenfassung aller Ordinalzahlen zu einem Ganzen führt also zu Widersprüchen. Ω ist, wie die Russell-Zermelo-Klasse R oder das Universum V , zu groß, um eine Menge zu sein. Die Reihe der Ordinalzahlen

$$0, 1, 2, \dots, \omega, \dots, \alpha, \dots$$

ist, im Cantorschen Sinne, absolut unendlich.

Man kann auch wie folgt argumentieren:

Annahme, Ω ist eine Menge. Sei dann $\delta = \sup(\Omega)$. Dann ist $\delta \in \Omega$.

Da keine größte Ordinalzahl existiert, ist das Supremum echt, und daher gilt $\alpha < \delta$ für alle $\alpha \in \Omega$, insbesondere also $\delta < \delta$, *Widerspruch!*

Ebenso zeigt man, daß die Kardinalzahlen keine Menge bilden:

Alephparadoxon von Cantor

Wir setzen

$$\aleph = \{ \aleph_\alpha \mid \alpha \in \Omega \}.$$

[$\aleph$ = taw ist der letzte Buchstabe des hebräischen Alphabets,

wie Ω der letzte Buchstabe des griechischen Alphabets ist.

Die Bezeichnung verwendet Cantor in einem Brief an Dedekind, s. u.]

Annahme, $\aleph$ ist eine Menge. Sei dann $\aleph^* = \sup(\aleph)$.

Da das Supremum einer Menge von Kardinalzahlen wieder eine Kardinalzahl ist, gilt $\aleph^* \in \aleph$. Andererseits hat $\aleph$ kein größtes Element, denn ist $\kappa \in \aleph$, $\kappa = \aleph_\alpha$, so ist $\kappa < \aleph_{\alpha+1}$.

Also ist das Supremum echt, und damit gilt $\kappa < \aleph^*$ für alle $\kappa \in \aleph$.

– Insbesondere also $\aleph^* < \aleph^*$. *Widerspruch!*

Nach dem Wohlordnungssatz ist $\aleph = \text{Kard} = \{ \alpha \mid \alpha \text{ ist Kardinalzahl} \}$. Dies wird für das Argument nicht gebraucht. Wir brauchen, daß κ^+ existiert für alle $\kappa \in \aleph$, d. h. wir brauchen: Für alle Ordinalzahlen α existiert eine Ordinalzahl β mit $|W(\alpha)| < |W(\beta)|$. Der Satz von Hartogs liefert uns diese Aussage (oder natürlich der Wohlordnungssatz). Cantor kannte in den Jahren vor 1900, als er sich intensiv mit dem Phänomen der inkonsistenten Zusammenfassungen beschäftigte, den Satz von Hartogs noch nicht, und den Wohlordnungssatz wollte er gerade beweisen. Er rechtfertigt die Existenz von κ^+ etwa so: Sei $\aleph_\alpha \in \aleph$, und sei $Z_\alpha = \{ \beta \mid \beta \text{ ist Ordinalzahl und } |W(\beta)| \leq \aleph_\alpha \}$. Dann ist $\langle Z_\alpha, < \rangle$ eine Wohlordnung. Sei also $\gamma = \text{o. t.}(\langle Z_\alpha, < \rangle)$. Offenbar gilt dann einfach $Z_\alpha = W(\gamma)$, und also insbesondere $\gamma \notin Z_\alpha$. Also $\aleph_\alpha < |W(\gamma)|$. Dann aber $|W(\gamma)| = \aleph_{\alpha+1}$, denn für jedes $\beta < \gamma$ gilt $|W(\beta)| \leq \aleph_\alpha$.

Der Punkt hier ist: Warum dürfen wir Z_α bilden? Oder genauer: Warum ist die Zusammenfassung Z_α eine konsistente Vielheit? Nun, wir haben auch $\mathcal{P}(M)$ gebildet für viele Mengen M . Man wird langfristig um eine axiomatische Einstellung, die derlei Existenzen sichert, nicht herumkommen. Anders ausgedrückt: Man fragt zurück: Warum nicht?

Cantor ist in der Tat selbst auf dieses Problem eingegangen:

Cantor (28. 8. 1899, Brief an Dedekind): „Man muß die Frage aufwerfen, woher ich denn wisse, daß die wohlgeordneten Vielheiten oder Folgen, denen ich die Kardinalzahlen

$$\aleph_0, \aleph_1, \dots, \aleph_{\omega_0}, \dots, \aleph_{\omega_1}, \dots$$

zuschreibe, auch wirklich ‚Mengen‘ in dem erklärten Sinne des Wortes, d. h. ‚konsistente Vielheiten‘ seien. Wäre es nicht denkbar, daß schon *diese* Vielheiten ‚inkonsi-

stent‘ seien, und daß der Widerspruch der Annahme eines ‚Zusammenseins aller ihrer Elemente‘ sich *nur noch nicht bemerkbar gemacht hätte?* Meine Antwort hierauf ist, daß diese Frage auf *endliche Vielheiten ebenfalls auszudehnen* ist und daß eine genaue Erwägung zu dem Resultate führt: sogar für endliche Vielheiten ist ein ‚Beweis‘ für ihre ‚Konsistenz‘ nicht zu führen. Mit anderen Worten: Die Tatsache der ‚Konsistenz‘ endlicher Vielheiten ist eine einfache unbeweisbare Wahrheit, es ist ‚*Das Axiom* der Arithmetik (im alten Sinne des Wortes)‘. Und ebenso ist die ‚Konsistenz‘ der Vielheiten, denen ich die Alephs als Kardinalzahlen zuspreche, ‚das Axiom der erweiterten, der transfiniten Arithmetik.‘

Ich würde sehr gerne diese Sachen mündlich mit Ihnen genauer durchsprechen ... Allein ich weiß nicht, ob es nicht für Sie momentan eine Störung in Ihren Arbeiten wäre, wenn ich mich von hier aus für ein paar Stunden zu Ihnen auf den Weg machte?“

Was sich da kurz vor der Jahrhundertwende skizzenhaft abzeichnet, ist seiner Zeit so weit voraus, daß es kaum verwundert, daß ein mündlicher Austausch mit Zeitgenossen schnell im Sand verlief. Die Zweifel an „Konsistenz“-Beweisen für den endlichen Bereich der Mengenwelt sind bemerkenswert. Und für die transfinite Mengenlehre selber ist es weit mehr als eine Axiomatik, die Cantor vorschwebt. *Das Axiom* der transfiniten Arithmetik ist ein Metaaxiom, oder, wenn man so will, ein Manifest. Es identifiziert Konsistenz und Existenz von Ordinalzahlkomprehensionen. Erst die Inkonsistenz aller Alephs oder aller Ordinalzahlen ist die obere Grenze, das Ende des Weges. Davor ist alles möglich, und solange der Weg der Ordinalzahlen nicht aus Konsistenzgründen den sicheren Boden unter den Füßen verliert, garantiert *das Axiom* die Existenz fraglicher Zusammenfassungen, im einfachsten Fall etwa die Existenz von Nachfolgerkardinalzahlen. Dies ist ein weit verbreitetes Bild auch der heutigen Mengenlehre, und es schließt mühelos die Existenz großer Kardinalzahlen mit ein. Die Folge der Ordinalzahlen ist so lang und so reich, wie es nur irgendwie geht. Sie selbst ist ihre eigene Grenze und ihr eigenes Ziel. Haben wir α , so haben wir $\alpha + 1$. Weiter bilden wir Limeszahlen und Limeszahlen von Limeszahlen. Und wir lassen den Prozeß ab α solange laufen, bis wir ein β erreichen derart, daß $W(\beta)$ selbst nicht mehr die Größe von $W(\alpha)$ hat. ‚Fertig‘-Machen. Weiterzählen. Iterieren. Alles Bisherige nehmen, und damit etwas Neues in Gang setzen. (Für Neumann-Zermelo-Ordinalzahlen ist *Alles Bisherige* zugleich schon das Neue selber.) Erst, wenn *Alles Bisherige* wirklich *Alles* ist, sind wir fertig und haben die Ordinalzahlen vor uns ausgerollt, gesetzt, daß wir jemals so weit kommen. Andernfalls können wir nur abstrahieren und den Prozeß lediglich als Ganzes von außen betrachten und nicht mit Worten von innen ausleuchten, und dann ist die Inkonsistenz der Menge aller Ordinalzahlen alles, was wir vom Ende des Weges wissen. Cantor hat, mit metaphysisch-religiösen Motiven, in dieser Weise gedacht, bei der die Paradoxien zum festen Bestandteil der Mengenlehre werden, sie nicht bedrohen, sondern ihr einen einzigen, absolut unendlich fernen Punkt vorenthalten, unterhalb dessen sie sich frei entfaltet. Die Erkenntnis der letzten Grenze ist schließlich menschlicher Triumph, und ein nobles Wissen:

Cantor (2. 10. 1897, *Brief an Hilbert*): „Sie übersehen jedoch, daß ich in meinem Harzburger Schreiben noch das Charakteristikum ‚fertig‘ gebraucht und gesagt habe:

Theorem: ‚Die Totalität aller Alephs läßt sich nicht als bestimmte und zugleich fertige Menge auffassen.‘

Hierin ist das *punctum saliens* zu sehen und ich wage es, dieses *vollkommen sichere*, aus der *Definition der ‚Totalität aller Alephs‘ streng beweisbare Theorem* als den wichtigsten und vornehmsten Satz der Mengenlehre zu bezeichnen.“

Die Mengenlehre wird zum Gentleman. Sie kennt ihre Möglichkeiten und Grenzen, und nutzt und akzeptiert beide. Man muß sich einmal vor Augen führen, was Cantor aus dem Begriff der linearen Ordnungen, in denen nichtleere Teilmengen ein kleinstes Element haben, gemacht hat. Unter seiner Hand wird dieses Rohmaterial zum grandiosen Weg in einen metaphysischen Nebel. Dieser Weg verläßt nie das sichere Land der Mathematik und bringt doch so viel an Gedankenreichtum mit sich, daß die Mengenlehre zu einer wahrhaft menschlichen Tätigkeit wird. Cantor, der über Herbarts wandelbare Grenze polemisiert hat, steigt nun selber frei die absolut gegebenen Ordinalzahlen empor, *das Axiom* garantiert ihm, daß er dabei stets weit vom großen Feuerwerk des Endes entfernt ist. Nach ihm ist es erst wieder Paul Mahlo, der ähnlich frei ist, und der für die Mitwelt exorbitant große Kardinalzahlen mit einer Gelassenheit betrachtet, die seine Arbeiten zu einem singulären Ereignis in der Geschichte der Mengenlehre werden ließ – auch wenn Mahlos Kardinalzahlen heute eher schon wieder als klein empfunden werden.

Cantor hat die inkonsistente Vielheit aller Ordinalzahlen als Argument für die Wohlordenbarkeit jeder Menge verwendet: Wir tragen eine Menge M rekursiv ab. Würden wir dabei niemals zu Ende kommen, so entstünde eine Bijektion zwischen Ω und einer Teilmenge N von M . Mit M ist aber auch N eine Menge. Aber N ist dann gleichmächtig mit Ω , also wäre auch Ω eine Menge. Hier finden zwei „Axiome“ über konsistente Vielheiten Anwendung: Teilvielheiten von Mengen sind Mengen, und ist von zwei bijektiv aufeinander abbildbaren Vielheiten die eine eine Menge, so ist es auch die andere. (Cantor hat Hilbert 1898 vier Axiome mitgeteilt, die der Leser am Ende von Abschnitt 3, Kapitel 1 finden wird.)

Briefe von Cantor an Hilbert und Dedekind schildern diesen Beweisversuch des Wohlordnungssatzes. Cantors Ziel war es zu zeigen, daß die Kardinalität einer beliebigen Menge ein Aleph ist: Für Cantor ist die Kardinalität einer Menge immer definiert durch einen zweifachen Abstraktionsprozeß; die Alephs dagegen sind spezielle Kardinalitäten, nämlich die Kardinalzahlen, die den wohlgeordneten Mengen zukommen. In unserer Terminologie bedeutet Cantors Ausdruck „jede Kardinalzahl ist ein Aleph“ einfach „jede Menge läßt sich bijektiv auf ein $W(\beta)$ abbilden für eine Ordinalzahl β “. Es geht also um den Wohlordnungs- und den Vergleichbarkeitssatz.

Cantor (26. 9. 1897, *Brief an Hilbert*): „Wenn eine bestimmte wohldefinierte *fertige* Menge eine Kardinalzahl haben würde, die mit keinem der Alephs zusammenfiel, so müßte sie Teilmengen enthalten, deren Kardinalzahl *irgend* ein Aleph ist, oder mit anderen Worten, die Menge müßte die Totalität aller Alephs in sich tragen.‘

Daraus ist leicht zu folgern, daß unter der gegebenen Voraussetzung (einer *bestimmten Menge*, deren Kardinalzahl kein Aleph wäre) auch die Totalität aller Alephs als eine bestimmte wohldefinierte fertige Menge aufgefaßt werden könnte. Daß dies nicht der Fall ist, habe ich oben bewiesen. Es ist daher jedes α auch immer ein bestimmtes Aleph.“

Cantor (3. 8. 1899, *Brief an Dedekind*): „Es erhebt sich nun die Frage, ob in diesem System $\aleph$ [aller Alephs $\aleph_0, \aleph_1, \dots, \aleph_\omega, \aleph_{\omega+1}, \dots$] *alle transfiniten Kardinalzahlen* enthalten sind? Gibt es, mit anderen Worten, eine *Menge*, deren Mächtigkeit kein Aleph ist?

Diese Frage ist zu *verneinen*, und der Grund dafür liegt in der von uns erkannten Inkonsistenz der Systeme Ω [aller Ordinalzahlen] und $\aleph$.

Beweis: Nehmen wir eine bestimmte Vielheit V und setzen voraus, daß ihr *kein Aleph als Kardinalzahl* zukommt, so schließen wir, daß V *inkonsistent* sein muß.

Denn man erkennt leicht, daß unter der gemachten Voraussetzung das ganze System Ω in die Vielheit hineinprojizierbar ist, d. h. daß eine Teilvielheit V' von V existieren muß, die dem System Ω äquivalent ist.

V' ist *inkonsistent*, weil Ω es ist, es muß also dasselbe auch von V behauptet werden.

Mithin muß jede transfinite *konsistente Vielheit*, jede transfinite Menge ein *bestimmtes Aleph* als Kardinalzahl haben, Also:

C. *Das System $\aleph$ aller Alephs ist nichts anderes als das System aller transfiniten Kardinalzahlen.*

... Wir erkennen ferner aus C die Richtigkeit des in [Cantor, 1895] ausgesprochenen [Vergleichbarkeits-] Satzes:

„Sind α und β beliebige Kardinalzahlen, so ist entweder $\alpha = \beta$ oder $\alpha < \beta$ oder $\alpha > \beta$.“

Denn die *Alephs* haben, wie wir gesehen, diesen Größencharakter.

So viel *in Kürze* für heute ... “

Cantors Intuition ist von der Kraft eines Herkules. Man kann über den Genauigkeitsgrad der Argumentation streiten, aber wie man das Argument formal sauber hinbekommt, steht auf einem ganz anderen Blatt. Unbestritten bleibt, daß Cantor, wo andere sich fürchteten, den Nemeischen Löwen der inkonsistenten Vielheiten mit der Keule erschlagen hat, und dann auch noch sein Fell nutzbringend verwendete.

Im dritten Abschnitt werden wir nun ein Axiomensystem angeben, das die Existenz gewisser Mengen fordert, und dessen Prinzipien wir in den meisten Fällen als die Erlaubnis lesen können, bestimmte Vielheiten zu Mengen zusammenzufassen. Das System erlaubt uns dagegen nicht die volle Mengenkompensation, und wir können in ihm V, Ω oder $\aleph$ nicht bilden. Während auf diese Weise die aufgetretenen Paradoxien wegfallen, ist das System andererseits reichhaltig genug, sodaß wir alle Objekte, die wir für unsere Sätze und ihre Beweise benötigen, darin als Mengen vorfinden.

Damit verlassen wir die Cantorsche Mengenlehre, bis wir sie in der axiomatischen Robe wiederfinden. Wir haben gesehen, daß Cantors Mengendefinition, die eine Menge als Zusammenfassung bestimmter Objekte zu einem Ganzen beschreibt, sehr sorgfältig gelesen werden muß. Die aufgetretenen Widersprüche

sind aber Widersprüche nur im Hinblick auf das naive Komprehensionsschema, das kein Teil der Mengenlehre Cantors ist. In seiner Mengenlehre sind die Paradoxien Erkenntnisse über zu große Zusammenfassungen, über absolut unendliche oder inkonsistente Vielheiten. Es ist nicht der Fall, daß mit der Eroberung des Unendlichen an allen Ecken und Enden das Chaos hereingebrochen ist. Es herrscht erst dort, so unsere Deutung der Paradoxien, wo wir einen unbeschränkten Teil des Mengenuniversums zu einem begrenzten machen wollen. So wie wir es in der Realität nur mit endlichen, nie mit aktual unendlichen Objekten zu tun haben, so kennt das Cantorsche, platonische Universum der Mathematik nur Mengen als Objekte, und keine absolut unendlichen Vielheiten.

Felix Hausdorff über zu große Zusammenfassungen

„Die Mengenlehre ist das Fundament der gesamten Mathematik; Differential- und Integralrechnung, Analysis und Geometrie arbeiten in Wirklichkeit, wenn auch vielleicht in verschleiender Ausdrucksweise, beständig mit unendlichen Mengen. Über das Fundament dieses Fundaments, also über eine einwandfreie Grundlegung der Mengenlehre selbst ist eine vollkommene Einigung noch nicht erzielt worden. Die nächstliegenden Schwierigkeiten und Vorurteile dürfen zwar als erledigt gelten: viele anscheinende ‚Paradoxien des Unendlichen‘ sind nur so lange paradox, wie man an der unberechtigten Forderung festhält, daß für endliche und für unendliche Mengen unterschiedslos dieselben Gesetze gelten sollen. Die naturgemäßen Abweichungen zwischen beiden Gebieten bedingen keinen Widerspruch innerhalb des Unendlichen. Dagegen ist eine wirkliche Paradoxie noch nicht befriedigend aufgeklärt, auf die der naive Mengenbegriff, mit seiner Zusammenfassung beliebig vieler Elemente zu einer Menge, letzten Endes hinausführt. In einer typischen Form, die allerdings noch der mathematischen Bestimmtheit entbehrt, lautet dies Paradoxon folgendermaßen: wenn es zu jeder Menge von Dingen noch ein weiteres, von ihnen allen verschiedenes Ding gibt, so ist die Gesamtheit aller Dinge offenbar selbst keine Menge. Ein solches System von Dingen, das nicht als Menge aufgefaßt werden kann, bilden, wie wir sehen werden, die (endlichen und unendlichen) Kardinalzahlen oder auch die Ordnungszahlen. Den hier nach notwendigen Versuch, den Prozeß der uferlosen Mengenbildung durch geeignete Forderungen einzuschränken, hat E. Zermelo unternommen. Da indessen diese äußerst scharfsinnigen Untersuchungen noch nicht als abgeschlossen gelten können und da eine Einführung des Anfängers in die Mengenlehre auf diesem Wege mit großen Schwierigkeiten verbunden sein dürfte, so wollen wir hier den naiven Mengenbegriff zulassen, dabei aber tatsächlich die Beschränkungen innehalten, die den Weg zu jenem Paradoxon abschneiden.“

(Felix Hausdorff 1914, „Grundzüge der Mengenlehre“)

*

*

*

1. Das Axiomensystem ZFC

Wir stellen in diesem Abschnitt das heute am häufigsten verwendete Axiomensystem ZFC für die Mengenlehre vor. „ZFC“ ist hierbei eine Abkürzung für „Zermelo-Fraenkel-Axiomatik mit Auswahlaxiom“. Das „C“ stammt hierbei von der englischen Bezeichnung „axiom of choice“ des Auswahlaxioms, das ein ausgezeichnetes Axiom der Zermeloschen Axiomatik darstellt.

Zermelo trug mit seiner Arbeit „Untersuchungen über die Grundlagen der Mengenlehre“, erschienen 1908 in den „Mathematischen Annalen“, den Hauptteil zu diesem System bei. Seine dort vorgestellten sieben Axiome wurden später nur noch ergänzt durch das Ersetzungsschema von Abraham Fraenkel (1922) und das Fundierungsaxiom, das auf Abraham Fraenkel (1922), John von Neumann (1925) und Ernst Zermelo (1930) zurückgeht. Thoralf Skolem (1887 – 1963) hat ebenfalls dem Ersetzungsschema und dem Fundierungsaxiom verwandte Prinzipien betrachtet, und zudem den Weg gewiesen, die axiomatische Mengenlehre durch die Verwendung einer exakten Sprache zu präzisieren (1922). In Briefen von Cantor an Hilbert finden sich axiomatische Ansätze, insbesondere hat Cantor bereits 1898 das Ersetzungsschema formuliert.

Die Abkürzung ZFC ist aus historischer Sicht unglücklich. Zermelos System enthält das Auswahlaxiom bereits als fundamentalen Bestandteil, und die Isolierung dieses Axioms stellt gerade eine der großen Leistungen Zermelos dar. Aus mathematischer Sicht ist die Betonung des „C“ in der Bezeichnung ZFC aber gerechtfertigt, da das Auswahlaxiom eine Sonderstellung unter den Axiomen einnimmt, und seine Verwendung im systematischen Aufbau der Mengenlehre und der restlichen Mathematik eine besondere Beachtung verdient.

Zermelo formulierte seine Axiomatik noch nicht in der Sprache der Prädikatenlogik erster Stufe, die heute die „offizielle“ Sprache der Mengenlehre ist. Auch wir werden zuerst die Axiome und einige elementare Folgerungen in der üblichen mathematischen Umgangssprache angeben. Der kritischen Richtung tragen wir dann im zweiten Kapitel vollends Rechnung, wo wir die Prädikatenlogik einführen und insbesondere den Begriff einer Eigenschaft präzisieren. Dort schaffen wir den formalen Rahmen für die Mengenlehre, der für metamathematische Untersuchungen und Resultate wie etwa dem der Unabhängigkeit der Kontinuumshypothese notwendig ist.

Im zweiten Band werden wir die Mengenlehre auf der Basis der ZFC-Axiome systematisch entwickeln. Die hier erzielten Ergebnisse der Zeit Cantors fügen sich in diesen axiomatischen Aufbau zwanglos ein. Wir können sie und ihre Beweise unverändert übernehmen, lediglich eine etwas andere Organisation des Stoffes ist notwendig, da wir z. B. keine Zahlen mehr als Grundobjekte zur Verfügung haben, sondern sie erst als Mengen konstruieren müssen.

Eine der Hauptaufgaben der mengentheoretischen Axiomatik, unabhängig von metamathematischen Gesichtspunkten, war es, die Theorie des Unendlichen zu retten, und dabei gleichzeitig die aufgetretenen Paradoxien zu eliminieren. Die Paradoxien verschwinden nun in natürlicher Weise in ZFC, da die Axiome keine uferlosen Zusammenfassungen zulassen, wie etwa die Bildung von $\{x \mid x \text{ ist Menge}\}$ oder allgemein die von $\{x \mid \mathcal{E}(x)\}$.

Zermelo schreibt in der Einleitung zu seiner Axiomatik von 1908:

Zermelo (1908b): „Die Mengenlehre ist derjenige Zweig der Mathematik, dem die Aufgabe zufällt, die Grundbegriffe der Zahl, der Anordnung und der Funktion in ihrer ursprünglichen Einfachheit mathematisch zu untersuchen und damit die logischen Grundlagen der gesamten Arithmetik und Analysis zu entwickeln; sie bildet somit einen unentbehrlichen Bestandteil der mathematischen Wissenschaft. Nun scheint aber gegenwärtig gerade diese Disziplin in ihrer ganzen Existenz bedroht durch gewisse Widersprüche oder ‚Antinomien‘, die sich aus ihren scheinbar denknotwendig gegebenen Prinzipien herleiten lassen und bisher noch keine allseitig befriedigende Lösung gefunden haben. Angesichts namentlich der ‚Russellschen Antinomie‘ von der ‚Menge aller Mengen, welche sich selbst nicht als Element enthalten‘ scheint es heute nicht mehr zulässig, einem beliebigen logisch definierbaren Begriff eine ‚Menge‘ oder ‚Klasse‘ als seinen ‚Umfang‘ zuzuweisen. Die ursprüngliche Cantorsche Definition einer ‚Menge‘ als einer ‚Zusammenfassung von bestimmten wohlunterschiedenen Objekten unserer Anschauung oder unseres Denkens zu einem Ganzen‘ bedarf also jedenfalls einer Einschränkung, ohne daß es doch schon gelungen wäre, sie durch eine andere, ebenso einfache zu ersetzen, welche zu keinen solchen Bedenken mehr Anlaß gäbe. Unter diesen Umständen bleibt gegenwärtig nichts anderes übrig, als den umgekehrten Weg einzuschlagen und, ausgehend von der historisch bestehenden ‚Mengenlehre‘, die Prinzipien aufzusuchen, welche zur Begründung dieser mathematischen Disziplin erforderlich sind. Diese Aufgabe muß in der Weise gelöst werden, daß man die Prinzipien einmal eng genug einschränkt, um alle Widersprüche auszuschließen, gleichzeitig aber auch weit genug ausdehnt, um alles Wertvolle dieser Lehre beizubehalten.“

In der hier vorliegenden Arbeit gedenke ich nun zu zeigen, wie sich die gesamte von G. Cantor und R. Dedekind geschaffene Theorie auf einige wenige Definitionen und auf sieben anscheinend voneinander unabhängige ‚Prinzipien‘ oder ‚Axiome‘ zurückführen läßt. Die weitere, mehr philosophische Frage nach dem Ursprung und dem Gültigkeitsbereiche dieser Prinzipien soll hier noch unerörtert bleiben. Selbst die gewiß sehr wesentliche ‚Widerspruchslosigkeit‘ meiner Axiome habe ich noch nicht streng beweisen können, sondern mich auf den gelegentlichen Hinweis beschränken müssen, daß die bisher bekannten ‚Antinomien‘ sämtlich verschwinden, wenn man die hier vorgeschlagenen Prinzipien zugrunde legt. Für spätere Untersuchungen, welche sich mit solchen tiefer liegenden Problemen beschäftigen, möchte ich hiermit wenigstens eine nützliche Vorarbeit liefern.

Der nachstehende Artikel enthält die Axiome und ihre nächsten Folgerungen ...“

Das Problem der Widerspruchsfreiheit ist systemimmanent und nicht behebbar. Nach dem zweiten Gödelschen Unvollständigkeitssatz kann die Widerspruchsfreiheit von ZFC nicht innerhalb von ZFC bewiesen werden (es sei denn, ZFC ist widerspruchsvoll, denn dann ist alles beweisbar). Aus der Universalität

von ZFC – der Tatsache, daß jedes mathematische Argument in ZFC ausgeführt werden kann – folgt, daß wir den Zusatz „innerhalb von ZFC“ streichen können. Wir erhalten dann: Die Widerspruchsfreiheit von ZFC ist nicht beweisbar (es sei denn, ZFC ist widerspruchsvoll). Daß bei der Untersuchung einer reichhaltigen Axiomatik keine Widersprüche auftreten und Widersprüche anderer Systeme verschwinden, ist streng genommen alles, was wir zur Widerspruchsfreiheit einer solchen Axiomatik sagen können. Aus humanistischer Sicht spricht aber sicher ein kohärentes und ästhetisches Gesamtbild einer mathematischen Theorie für ihre Widerspruchsfreiheit, und in dieser Hinsicht sind ZFC und die modernen Erweiterungen von ZFC sicherlich widerspruchsfrei.

Die Bedeutung von „x existiert“

Eine Welt für die Mengenlehre, ein Modell, in welchem die Axiome der Mengenlehre gelten, nennt Zermelo einen Bereich. Ein Bereich besteht aus Dingen. Die Dinge eines Zermeloschen Bereiches zerfallen in Mengen und Urelemente (Grundobjekte). Der intendierte, durch die Axiome zu beschreibende Bereich ist das platonische Mengenuniversum, aber man wird bestenfalls hoffen können, daß die Axiome für genau einen Bereich zutreffen werden. Auf diesen Punkt der verschiedenen Modelle für die Mengenlehre kommen wir im zweiten Band ausführlich zurück.

Zermelo (1908b):

„§ 1 Grundlegende Definitionen und Axiome

1. Die Mengenlehre hat zu tun mit einem ‚Bereich‘ $\mathfrak{B}$ von Objekten, die wir einfach als ‚Dinge‘ bezeichnen wollen, unter denen die ‚Mengen‘ einen Teil bilden. Solen zwei Symbole a und b dasselbe Ding bezeichnen, so schreiben wir $a = b$, im entgegengesetzten Falle $a \neq b$. Von einem Dinge a sagen wir, es ‚existiere‘, wenn es dem Bereiche $\mathfrak{B}$ angehört ...

2. Zwischen den Dingen des Bereiches bestehen gewisse ‚Grundbeziehungen‘ der Form $a \in b$. Gilt für zwei Dinge a, b die Beziehung $a \in b$, so sagen wir, ‚ a sei *Element* der Menge b ‘ oder ‚ b enthalte a als Element‘ oder ‚besitze das Element a ‘. Ein Ding b , welches ein anderes a als Element enthält, kann immer als eine *Menge* bezeichnet werden, aber auch nur dann – mit einer einzigen Ausnahme (Axiom II) [die Ausnahme ist die leere Menge].

3. ... [Teilmengen und disjunkte Mengen]

4. Eine Frage oder Aussage $\mathfrak{G}$, über deren Gültigkeit oder Ungültigkeit die Grundbeziehungen des Bereiches vermöge der Axiome und der allgemeingültigen logischen Gesetze ohne Willkür entscheiden, heißt ‚*definit*‘. Ebenso wird auch eine Klassenaussage $\mathfrak{G}(x)$, in welcher der variable Term x alle Individuen einer Klasse $\mathfrak{K}$ durchlaufen kann, als ‚definit‘ bezeichnet, wenn sie für *jedes einzelne* Individuum x der Klasse $\mathfrak{K}$ definit ist. So ist die Frage, ob $a \in b$ oder nicht ist, immer definit, ebenso die Frage, ob $M \subseteq N$ oder nicht.

Über die Grundbeziehungen unseres Bereiches $\mathfrak{B}$ gelten nun die folgenden ‚*Axiome*‘ oder ‚*Postulate*‘ ...“

Gemeint ist mit „definiter Aussage“ eine „mathematische Aussage über Mengen“. Eine präzise Definition geben wir erst im nächsten Kapitel, für jetzt begnügen wir uns, wie Zermelo in seiner Arbeit, mit einem hohen, aber nicht maximalen Grad an Exaktheit.

Der Zermelosche Begriff des Bereiches hat – wie der Begriff *definit* – Kritik hervorgerufen. Er ist aber einfach eine andere Bezeichnung für das, was wir bisher Mengenuniversum nannten – bei einem gleichzeitigen Verzicht, von *einem* Mengenuniversum schlechthin zu reden. Bezweifelt man, daß eine *Existenzaussage* über Mengen (oder mathematische Objekte) irgendetwas mit der *Existenz* eines Objektes zu tun hat, so wird man diese Kritik sicherlich teilen, und wahrscheinlich einen formalistischen Standpunkt vertreten. Für diesen Standpunkt, auf den wir im zweiten Kapitel noch genauer zurückkommen, ist eine Existenzaussage einfach eine Zeichenkette, und alles weitere ist eine mehr oder weniger fruchtbare Zutat der menschlichen Phantasie.

Für den Platoniker dagegen gibt es einen ausgezeichneten Bereich, über den die Axiome reden, nämlich die Welt aller Mengen, das mengentheoretische Universum, das für sich und unabhängig von uns existiert – ganz so, wie die natürlichen Zahlen für viele Mathematiker absolut existieren, wenn auch jeder eine etwas andere Vorstellung von ihnen haben mag. Die Axiome der Mengenlehre konstatieren einfach einige richtige Aussagen über das mengentheoretische Universum. Daß sie auch noch innerhalb anderer Modelle gelten können, innerhalb von Teilbereichen des Universums, spiegelt lediglich den Reichtum der platonischen Welt aller Mengen wider, und die Begrenztheit eines Axiomensystems, das versucht, diesen Reichtum vollständig zu erfassen.

Zwischen diesen beiden Extremen der arktischen Trostlosigkeit und des süditalienischen Katholizismus gibt es eine Vielzahl von möglichen Positionen, und es gibt wenige an Grundlagenfragen interessierte Mathematiker, die Zeit ihres Lebens auf der Weltkugel der Interpretation immer am gleichen Ort residieren.

Die einzelnen Axiome

Wir besprechen nun die Axiome von Zermelo, Fraenkel und von Neumann. Diese Axiome gelten in einem Bereich aus Objekten und Mengen, einer Welt aus „Dingen“. Wir lassen ab hier – im Gegensatz zur Einführung und zu Zermelo – keine Grundobjekte oder Urelemente mehr zu, die Zahlen $\mathbb{N}$, $\mathbb{Z}$, $\mathbb{Q}$ und $\mathbb{R}$ lassen sich leicht als Mengen einführen (etwa $\mathbb{N}$ als die Menge der endlichen Ordinalzahlen nach von Neumann, $\mathbb{Z}$ als $\langle \mathbb{N} - \{0\}, <^* \rangle + \langle \mathbb{N}, < \rangle$, $\mathbb{Q}$ als irgendeine geeignete Fortsetzung von $\langle \mathbb{Z}, < \rangle$ des Typs η , $\langle \mathbb{R}, < \rangle$ als die Dedekind-Vervollständigung von $\langle \mathbb{Q}, < \rangle$). Auch die Arithmetik auf den Zahlen ist leicht zu definieren und umständlich zu entwickeln, auf Neudeutsch *easy but tedious*). Da wir auch eine Grundlegung der ganzen Mathematik im Auge haben, ist der Verzicht auf Urelemente nur natürlich. Sie würden die Diskussion nur unnötig verkomplizieren. Je weniger man für eine Grundlegung der Mathematik als gegeben annimmt, desto besser. Wir halten also fest:

Alle Objekte sind Mengen, und die einzigen Grundrelationen zwischen Mengen sind die Gleichheit „ $a = b$ “ und die Elementbeziehung „ $a \in b$ “.

Auf eine Definition von „Menge“ und „ a ist Element von b “ wird bewußt verzichtet. Diese Grundbegriffe werden durch die Axiome lediglich beschrieben.

Fraenkel (1928): „Dieser Aufbau der Mengenlehre verläuft nach der sog. *axiomatischen Methode*, die vom historischen Bestand einer Wissenschaft (hier der Mengenlehre) ausgeht, um durch logische Analyse der darin enthaltenen Begriffe, Methoden und Beweise die zu ihrer Begründung erforderlichen Prinzipien – die *Axiome* – aufzusuchen und aus ihnen die Wissenschaft deduktiv herzuleiten. Gemäß dem Wesen dieser Methode sehen wir ganz davon ab, den Mengenbegriff zu *definieren* oder näher zu zergliedern; vielmehr gehen wir lediglich von gewissen Axiomen aus, in denen der Mengenbegriff wie auch die Relation ‚als Element enthalten sein‘ auftritt und die Existenz gewisser Mengen gefordert wird. Durch die Gesamtheit der Axiome wird so der Begriff der Menge gewissermaßen unausgesprochen festgelegt ...

Diese ganz formale, jedes sachlichen Inhalts entkleidete Auffassung, die auf eine Charakterisierung der Begriffe ‚Menge‘ und ‚Element sein‘ durch *Definition* bewußt verzichtet und sich mit einer Festlegung der Beziehungen zwischen den uns interessierenden Begriffen begnügt, wird vielleicht deutlicher durch den ... Hinweis, daß jede *mit den Axiomen verträgliche* Interpretation der Begriffe ‚Menge‘ und ‚als Element enthalten sein‘ zulässig und mit jeder anderen solchen Interpretation gleichberechtigt ist ... Man kann hiernach naturgemäß niemals entscheiden, was eine Menge ‚an sich‘ ist, ob z. B. ein Pferd eine Menge darstellt; eine inhaltliche Bestimmung der Grundbegriffe entspricht eben gar nicht dem Wesen der axiomatischen Methode. Vielmehr liegt eine nur implizite, unausgesprochene, überdies gegenseitig verkettete Definition der Begriffe ‚Menge‘ und ‚Elementsein‘ mittels der Gesamtheit der in den Axiomen auftretenden Aussagen vor.

Dieser Sachverhalt kann *cum grano salis* auch durch Vergleich mit anderen Wissenschaften, z. B. der Physik, verdeutlicht werden. Begriffe wie der der Wärme oder der Elektrizität, auch etwa der des Äthers, sind in der Naturbetrachtung zunächst implizit gegeben durch ihre Wirkungen und Verknüpfungen bei den experimentell feststellbaren Tatsachen; das Wesen dieser physikalischen Begriffe wird anschaulicher (und auch unabhängiger von dem Wechsel der wissenschaftlichen Theorien) durch die Beschreibung charakteristischer Wirkungen gekennzeichnet als durch eine rein begriffliche Definition dessen, was beim augenblicklichen Stand der Forschung unter dem ‚Wesen‘ der Wärme, der Elektrizität, des Äthers verstanden wird.“

Explizit halten wir fest, daß wir auch die vollständige Induktion und die rekursive Definition über die natürlichen Zahlen $\mathbb{N}$ – ein Teil des Bereiches – nicht als Beweisprinzip betrachten, sondern entsprechende Induktions- und Rekursionssätze beweisen müssen – und können. Dies geschieht im Prinzip ganz so wie in 2.7. Will man die axiomatische Theorie ähnlich aufziehen wie dort die nichtaxiomatische Mengenlehre, so ist es günstig, etwa für einen Beweis des Satzes von Cantor-Bernstein, vor einer allgemeinen Ordinalzahl-Induktion und -Rekursion die Induktion und Rekursion über $\mathbb{N}$ zu beweisen, was leicht geschehen kann, sobald man $\mathbb{N}$ selbst definiert hat.

Das Extensionalitätsaxiom

Das erste Axiom ist das bekannte Gleichheitsprinzip:

(EXT) Extensionalitätsaxiom

Zwei Mengen sind genau dann gleich, wenn sie die gleichen Elemente haben.

Das Axiom drückt aus, daß eine Menge vollständig durch ihre Elemente bestimmt ist. Das Denken der Mengenlehre ist durchweg extensional. Ein Begriff wird mit seinem Umfang identifiziert (vgl. die extensionale Definition einer Funktion als Menge von geordneten Paaren).

Wir definieren $a \subseteq b$ wie früher und erhalten aus dem Extensionalitätsaxiom: Für alle x, y gilt $x = y$ genau dann, wenn $x \subseteq y$ und $y \subseteq x$.

Zermelo (1908b):

„**Axiom I.** Ist jedes Element einer Menge M gleichzeitig Element von N und umgekehrt ... so ist immer $M = N$. Oder kürzer: jede Menge ist durch ihre Elemente bestimmt.

(Axiom der Bestimmtheit.)“

Drei einfache Existenzaxiome

Die weiteren Axiome sind nun bis auf das Fundierungsaxiom Existenzaxiome. Sie behaupten die Existenz bestimmter Mengen, wobei in den meisten Fällen Mengen mit Hilfe anderer Mengen gebildet werden. Die drei einfachsten Existenzaxiome betreffen die Existenz der Elementarmengen $\{a_1, \dots, a_n\}$ und der Vereinigungsmenge.

(LM) Existenz der leeren Menge

Es existiert eine Menge, welche keine Elemente hat.

Wir bezeichnen diese Menge wieder mit $\emptyset$ oder $\{ \}$. Nach dem Extensionalitätsaxiom ist die leere Menge eindeutig bestimmt: Es gibt genau eine Menge, die kein Element enthält.

(PA) Paarmengenaxiom

Zu je zwei Mengen x, y existiert eine Menge z , die genau x und y als Elemente hat.

Wir schreiben $\{x, y\}$ für die Paarmenge von x und y .

x und y müssen nicht verschieden sein, und es folgt für eine Menge x die Existenz von $\{x, x\} = \{x\}$, also die Existenz der Einermenge von x .

Das geordnete Paar (x, y) ist wieder definiert durch $(x, y) = \{\{x\}, \{x, y\}\}$.

Es gilt $(x_1, x_2) = (y_1, y_2)$ gdw $x_1 = y_1$ und $x_2 = y_2$.

Weiter ist z. B. $(x, x) = \{\{x\}\}$.

Relationen und Funktionen sowie Begriffe und Schreibweisen in diesem Umfeld sind wie im ersten Abschnitt definiert.

Zermelo (1908b):

„**Axiom II.** Es gibt eine (uneigentliche) Menge, die ‚Nullmenge‘ 0 , welche gar keine Elemente enthält. Ist a irgendein Ding des Bereiches, so existiert eine Menge $\{a\}$, welche a und nur a als Element enthält; sind a, b zwei Dinge des Bereiches, so existiert immer eine Menge $\{a, b\}$, welche sowohl a als [auch] b , aber kein von beiden verschiedenes Ding x als Element enthält.

(Axiom der Elementarmengen.)“

(VER) Vereinigungsmengenaxiom

Zu jeder Menge x existiert eine Menge y , deren Elemente genau die Elemente der Elemente von x sind.

Diese nach dem Extensionalitätsaxiom eindeutig bestimmte Vereinigungsmenge y von x bezeichnen wir mit $\bigcup x$. Es gilt dann für alle y :

$y \in \bigcup x$ gdw „es existiert ein $z \in x$ mit $y \in z$ “.

Ziehen wir das Paarmengenaxiom mit heran, so folgt, daß auch die Vereinigung von zwei Mengen existiert. Wir definieren hierzu für x und y :

$x \cup y = \bigcup \{x, y\}$.

Es gilt dann $z \in x \cup y$ gdw $z \in x$ oder $z \in y$.

Das Paarmengenaxiom garantiert uns die Existenz von $\{x, y\}$ für alle x, y . Für die Existenz der Dreiermenge $\{x, y, z\}$ usw. braucht man das Vereinigungsmengenaxiom. Wir setzen

$\{x, y, z\} = \bigcup \{\{x, y\}, \{z\}\},$

wobei wir hier dreimal das Paarmengenaxiom und einmal das Vereinigungsmengenaxiom bemühen. Allgemein definieren wir:

$$\begin{aligned}
\{x_1, x_2, x_3\} &= \bigcup \{ \{x_1, x_2\}, \{x_3\} \}, \\
\{x_1, x_2, x_3, x_4\} &= \bigcup \{ \{x_1, x_2, x_3\}, \{x_4\} \}, \\
&\dots \\
\{x_1, \dots, x_{n+1}\} &= \bigcup \{ \{x_1, x_2, x_3, \dots, x_n\}, \{x_{n+1}\} \}.
\end{aligned}$$

Wiederholte Anwendung des Paarmengen- und des Vereinigungsmengenaxioms zeigt dann für ein beliebiges n :

- (+) für jede Liste $x_1, \dots, x_n$ von Mengen existiert eine Menge $\{x_1, \dots, x_n\}$, die genau $x_1, \dots, x_n$ als Elemente enthält.

Nach (EXT) ist sie eindeutig bestimmt.

Die „metamathematischen“ natürlichen Zahlen $1, 2, \dots, n, n+1$ sind hier keine Objekte des Bereichs (Mengen), sondern lediglich Indikatoren für die Länge der Liste $x_1, \dots, x_n$. Zur Vermeidung von Verwechslungen mit Elementen des erst zu definierenden $\mathbb{N} = \omega$ des Objektbereichs ist es zuweilen suggestiv, sich die metamathematische Zahlenreihe $1, 2, 3, \dots$ als $1, 11, 111$, usw. vorzustellen. Üblicherweise werden aber die Zeichen $1, 2, 3, \dots$ sowohl zur Bezeichnung von Mengen verwendet, als auch zur Bezeichnung von metamathematischen natürlichen Zahlen wie oben bei der Definition von $\{x_1, \dots, x_n\}$. Der Leser muß dann entscheiden, auf welcher Ebene man sich befindet, was normalerweise problemlos ist.

Wir lassen auf der „Metaebene“ Induktion und Rekursion nach $1, 2, 3, \dots$ zu, und man kann etwa (+) durch Induktion beweisen. Aussagen wie (+) werden in der endlichen Welt, in der wir über Mathematik reden, nie wirklich für alle natürlichen Zahlen der Metaebene gebraucht. Wir haben irgendwann etwa mit Objekten zu tun, die wir mit $x_1, \dots, x_{10}$ bezeichnen, und zeigen dann durch mehrfache Anwendung von bestimmten Axiomen, daß $\{x_1, \dots, x_{10}\}$ existiert. Ein andermal mit 17 statt 10, oder mit 1220. Die Pünktchen „...“ wie in der Definition oben sind dann nur Hinweise, wie sich eine Definition für ein beliebiges, konkret gegebenes metamathematisches n , etwa 17, durchzuführen ist. Induktion ist aber eine bequeme Sprechweise in der Metaebene, und wir werden später etwa Aussagen über den Hilbert-Kalkül durch Induktion über die Länge von Beweisen zeigen. Wir brauchen aber nie wirklich einen Satz der Form „für alle $n \dots$ “, sondern immer nur ein Schema für beliebige, konkrete n .

Später noch einmal mehr zur Metaebene. Der Leser verwende Induktion und Rekursion über metamathematische natürliche Zahlen, woimmer es ihm angemessen erscheint. Oder er verwende Pünktchen „...“, und Sprechweisen wie „ n -malige Anwendung“, usw.

Analog sind Tripel, Quadrupel, $\dots$, n -Tupel für jede (metamathematische) natürliche Zahl n definiert durch:

$$\begin{aligned}
(x_1, x_2, x_3) &= ((x_1, x_2), x_3), \\
(x_1, x_2, x_3, x_4) &= ((x_1, x_2, x_3), x_4), \\
&\dots \\
(x_1, \dots, x_{n+1}) &= ((x_1, \dots, x_n), x_{n+1}).
\end{aligned}$$

Je zwei n -Tupel sind genau dann gleich, wenn sie in allen Komponenten übereinstimmen.

In ähnlicher Weise setzen wir

$$\begin{aligned}
x_1 \cup x_2 \cup x_3 &= (x_1 \cup x_2) \cup x_3, \\
x_1 \cup x_2 \cup x_3 \cup x_4 &= (x_1 \cup x_2 \cup x_3) \cup x_4, \\
&\dots \\
x_1 \cup \dots \cup x_{n+1} &= (x_1 \cup \dots \cup x_n) \cup x_{n+1}.
\end{aligned}$$

Damit gilt für alle $x_1, \dots, x_n$, daß

$$\{x_1, \dots, x_n\} = \{x_1\} \cup \{x_2\} \cup \dots \cup \{x_n\}.$$

Zermelo (1908b):

„**Axiom V.** Jeder Menge T entspricht eine Menge $\mathfrak{E}T$ (die ‚Vereinigungsmenge‘ von T), welche alle Elemente der Elemente von T und nur solche als Elemente enthält.

(Axiom der Vereinigung.)“

Das Aussonderungsschema

(AUS) Aussonderungsschema

Zu jeder Eigenschaft $\mathcal{E}$ und jeder Menge x gibt es eine Menge y , die genau die Elemente von x enthält, auf die $\mathcal{E}$ zutrifft.

Wir bezeichnen diese Menge y mit $\{u \in x \mid \mathcal{E}(u)\}$.

Der Name „Schema“ bezieht sich auf die frei wählbare Eigenschaft $\mathcal{E}$: Wir haben ein Axiom pro Eigenschaft, und unser Axiomensystem besteht damit aus unendlich vielen Axiomen. Wir werden im zweiten Band zeigen, daß sich ZFC nicht auf endlich viele Axiome reduzieren läßt.

Die Eigenschaft $\mathcal{E}$ darf andere Mengen als Parameter enthalten. Für jedes x können wir z. B. bilden:

$$\{u \in x \mid u \neq \emptyset\}.$$

Die verwendete Eigenschaft ist hier $\mathcal{E} = „u \neq \emptyset“$, mit der Menge $\emptyset$ als Parameter.

Der Unterschied zum vollen Komprehensionsschema ist, daß wir uns bei der Aufsammlung von Objekten mit der Eigenschaft $\mathcal{E}$ auf die Elemente einer fest vorgegebenen Menge x beziehen: Wir sondern aus x alle Elemente mit der Eigenschaft $\mathcal{E}$ aus, und fassen diese Elemente zu einer neuen Menge zusammen. Die entstehende Menge $\{u \in x \mid \mathcal{E}(u)\}$ ist immer eine Teilmenge von x , und damit ist das Aussonderungsschema ein „Existenzaxiom nach unten“. Obwohl in dieser Hinsicht das Aussonderungsschema viel schwächer ist als das uferlose Komprehensionsschema, können wir es in der Anwendung ähnlich flexibel einsetzen wie das Komprehensionsschema, da sich zumeist ein Teilbereich x des Mengenuniversums finden läßt, in dem alle gewünschten Objekte u mit der Eigenschaft $\mathcal{E}(u)$ als Elemente vorhanden sind. Ein Beispiel hierfür gibt das kartesische Produkt, das wir nach der Einführung des Potenzmengenaxioms definieren werden.

Das Aussonderungsschema garantiert uns die Existenz der Schnittmenge und der Subtraktion. Wir definieren:

$$x \cap y = \{z \in x \mid z \in y\},$$

$$x - y = \{z \in x \mid z \notin y\}.$$

Hierbei sind „ $z \in y$ “ und „ $z \notin y$ “ die verwendeten Eigenschaften. Die Menge y ist in beiden Fällen Parameter der Aussonderung aus x .

Übung

- (i) Definieren Sie $x_1 \cap \dots \cap x_n$.
- (ii) Definieren Sie $\emptyset$ mit Hilfe des Aussonderungsschemas und der Annahme der Existenz irgendeiner Menge.

Sei X eine nichtleere Menge und $x^* \in X$. Wir definieren:

$$\bigcap X = \{y \in x^* \mid \text{für alle } x \in X \text{ gilt } y \in x\}.$$

Offenbar ist $\bigcap X$ unabhängig von der Wahl von $x^* \in X$.

Zermelo (1908b):

„**Axiom III.** Ist die Klassenaussage $\mathfrak{G}(x)$ definit für alle Elemente einer Menge M , so besitzt M immer eine Untermenge $M_{\mathfrak{G}}$, welche alle diejenigen Elemente x von M , für welche $\mathfrak{G}(x)$ wahr ist, und nur solche als Elemente enthält.

(Axiom der Aussonderung.)“

Auswertung der Paradoxie von Russell-Zermelo

Bevor wir die weiteren Axiome von ZFC vorstellen, zeigen wir, daß sich die Konstruktion der Russell-Zermelo-Paradoxie nun in einen Satz verwandelt. Das Aussonderungsschema erlaubt uns die Konstruktion von $R = \{x \mid x \notin x\}$ nicht, und in diesem Sinne existiert die Paradoxie nicht mehr. Die Konstruktion liefert aber das folgende informative Ergebnis über das Mengenuniversum (oder, in der Sprache Zermelos, den zugrunde liegenden Bereich $\mathfrak{B}$):

Satz

Für alle x existiert ein $y \subseteq x$ mit $y \notin x$.

Beweis

Sei x beliebig. Nach dem Aussonderungsschema existiert

$$y = \{z \in x \mid z \notin z\}.$$

Dann gilt $y \notin x$, denn *andernfalls* hätten wir

$$y \in y \text{ gdw } y \notin y,$$

– wie im Paradoxon von Russell-Zermelo. *Widerspruch!*

Diese Überlegung bildet das erste *Theorem* in Zermelos Arbeit. Paradoxien werden, wie schon bei Cantor, zu mathematischen Resultaten.

Korollar

- (i) Es gibt kein x mit der Eigenschaft: Für alle y ist $y \in x$.
- (ii) Es gibt kein x mit der Eigenschaft: Für alle y mit $y \notin y$ ist $y \in x$.

Beweis

zu (i): Ein solches x würde alle seine Teilmengen als Elemente enthalten, im Widerspruch zum Satz oben.

zu (ii): Für ein solches x erhalten wir durch Aussonderung

$$x^* = \{z \in x \mid z \notin z\}.$$

Dann gilt nach Annahme über x für alle z : $z \in x^*$ gdw $z \notin z$.

– Für $z = x^*$ erhalten wir wieder $x^* \in x^*$ gdw $x^* \notin x^*$, *Widerspruch!*

(i) besagt: Das Mengenuniversum selbst ist keine Menge.

(ii) besagt: Es gibt viele Mengen x mit $x \notin x$. Denn obwohl wir den Fall $x \in x$ für eine Menge x nach den bisherigen Axiomen nicht ausschließen können, so sind doch die „normalen“ Mengen mit der Eigenschaft $x \notin x$ so zahlreich, daß sie in einer anderen Menge nicht alle enthalten sein können.

Nun aber zurück zu den Axiomen von ZFC.

Das Unendlichkeitsaxiom

(UN) Unendlichkeitsaxiom

Es existiert eine Menge x , die die leere Menge als Element enthält und die mit jedem ihrer Elemente y auch $\{y\}$ als Element enthält.

Es sind also die Mengen

$\emptyset, \{\emptyset\}, \{\{\emptyset\}\}, \{\{\{\emptyset\}\}\}, \dots$

Elemente von x . Eine kleinste Menge mit dieser Eigenschaft werden wir nach der Einführung des Potenzmengenaxioms besprechen.

Zermelo (1908b):

„**Axiom VII.** Der Bereich enthält mindestens eine Menge Z , welche die Nullmenge als Element enthält und so beschaffen ist, daß jedem ihrer Elemente a ein weiteres Element der Form $\{a\}$ entspricht, oder [anders formuliert,] welche mit jedem ihrer Elemente a auch die entsprechende Menge $\{a\}$ als Element enthält.

(Axiom des Unendlichen.)“

Ein in Verbindung mit dem Unendlichkeitsaxiom sehr starkes Existenzaxiom ist das Potenzmengenaxiom.

Das Potenzmengenaxiom

(POT) Potenzmengenaxiom

*Zu jeder Menge x existiert eine Menge y ,
die genau die Teilmengen von x als Elemente besitzt.*

Die Potenzmenge von x bezeichnen wir mit $\mathcal{P}(x)$.

Das Axiom garantiert, wie wir gesehen haben, die reiche Struktur der unendlichen Mengen hinsichtlich ihrer Mächtigkeit.

Das Potenzmengenaxiom kann man als eine zulässige Instanz des vollen Komprehensionsaxioms lesen:

Für jede Menge x existiert die Menge $\{y \mid \mathcal{E}(y)\}$, wobei $\mathcal{E}(y)$ die Eigenschaft „ $y \subseteq x$ “ ist.

Übung

Formulieren Sie die „Aufwärtsaxiome“ (LM), (PA), (VER) als Spezialfälle des vollen Komprehensionsschemas. [z. B. $\emptyset = \{x \mid x \neq x\}$.]

Den oben bewiesenen Satz, daß jede Menge eine Teilmenge besitzt, die nicht Element der Menge ist, können wir nun kurz so formulieren:

Satz

Es gibt kein x mit $\mathcal{P}(x) \subseteq x$.

Damit fällt auch die Mirimanovsche Paradoxie weg.

Das kartesische Produkt

Mit Hilfe des Potenzmengenaxioms und des Aussonderungsschemas können wir nun das kartesische Produkt definieren. Für Mengen A, B würden wir gerne zeigen, daß eine Menge $A \times B$ existiert mit der Eigenschaft: Für alle a, b gilt

$(a, b) \in A \times B$ gdw $a \in A$ und $b \in B$.

Die Definition $A \times B = \{(a, b) \mid a \in A \text{ und } b \in B\}$ ist unbeschränkte Komprehension und nicht durch das Aussonderungsschema abgesichert. Wir beobachten aber:

Für $a \in A, b \in B$ ist $(a, b) = \{\{a\}, \{a, b\}\} \subseteq \mathcal{P}(A \cup B)$.

Also $(a, b) \in \mathcal{P}(\mathcal{P}(A \cup B))$. Wir setzen also:

$A \times B = \{(a, b) \in \mathcal{P}(\mathcal{P}(A \cup B)) \mid a \in A \text{ und } b \in B\}$.

Dann gilt wie gewünscht $(a, b) \in A \times B$ gdw $a \in A$ und $b \in B$.

Weiter setzen wir:

$$\begin{aligned} A_1 \times A_2 \times A_3 &= (A_1 \times A_2) \times A_3, \\ A_1 \times A_2 \times A_3 \times A_4 &= (A_1 \times A_2 \times A_3) \times A_4, \\ &\dots \\ A_1 \times \dots \times A_{n+1} &= (A_1 \times \dots \times A_n) \times A_{n+1}. \end{aligned}$$

Für alle $A_1, \dots, A_n$ und alle x gilt dann:

$x \in A_1 \times \dots \times A_n$ gdw $x = (a_1, \dots, a_n)$ für gewisse $a_1 \in A_1, \dots, a_n \in A_n$.

Ist $A_1 = \dots = A_n$, so schreiben wir A^n für $A_1 \times \dots \times A_n$. A^n ist die Menge aller n -Tupel mit Einträgen aus A .

Zermelo (1908b):

„**Axiom IV.** Jeder Menge T entspricht eine zweite Menge $\mathcal{P}T$ (die ‚Potenzmenge‘ von T), welche alle Untermengen von T und nur solche als Elemente enthält.
(Axiom der Potenzmenge).“

Zermelos Zahlreihe Z_0

Eine Menge u heißt *Zermelo-induktiv* oder kurz *Z-induktiv*, falls gilt:

„ $\emptyset \in u$ und mit jedem $x \in u$ ist auch $\{x\} \in u$ “.

Eine Z-induktive Menge z^* existiert nach dem Unendlichkeitsaxiom. Wir setzen:

$$Z_0 = \bigcap \{ u \in \mathcal{P}(z^*) \mid u \text{ ist Z-induktiv} \}.$$

Diese Menge existiert nach dem Potenzmengenaxiom und dem Aussonderungsschema. Z_0 besteht intuitiv aus genau den Mengen $\emptyset, \{\emptyset\}, \{\{\emptyset\}\}, \dots$

Man zeigt leicht: Z_0 ist Z-induktiv. Weiter ist die Definition von Z_0 unabhängig von der speziellen Wahl der Z-induktiven Menge z^* (!).

Z_0 kann man als Menge der natürlichen Zahlen auffassen:

$$0 = \emptyset, \quad 1 = \{\emptyset\}, \quad 2 = \{\{\emptyset\}\}, \quad 3 = \{\{\{\emptyset\}\}\}, \quad \dots, \quad \text{also}$$

$$0 = \emptyset, \quad 1 = \{0\}, \quad 2 = \{1\}, \quad 3 = \{2\}, \dots$$

Zermelo (1908b): „... Die Menge Z_0 enthält die Elemente $0, \{0\}, \{\{0\}\}$ usw. und möge als ‚Zahlreihe‘ bezeichnet werden, weil ihre Elemente die Stelle der Zahlzeichen vertreten können. Sie bildet das einfachste Beispiel einer ‚abzählbar unendlichen‘ Menge ...“

Wir haben noch keine Induktion entlang einer Wohlordnung zur Verfügung, wenn wir über Z_0 etwas beweisen wollen. Wir müssen mit der Schnittdefinition auskommen. Diese läßt sich aber in einer induktiven Form notieren, denn direkt aus der Definition von Z_0 erhalten wir: Ist u eine Z-induktive Teilmenge von Z_0 , so ist $u = Z_0$. Dies liefert folgende Form der „Peano-Induktion“ für Z_0 :

„Sei $u \subseteq Z_0$ mit (1) $\emptyset \in u$, (2) für alle $x \in u$ ist $\{x\} \in u$. Dann ist $u = Z_0$.“ Hiermit zeigt man zum Beispiel:

Übung

- (i) Es gilt $x \neq \{x\}$ für alle $x \in Z_0$.

[Sei $u = \{x \in Z_0 \mid x \neq \{x\}\}$. Es genügt zu zeigen, daß u Z -induktiv ist. Offenbar gilt $\emptyset \in u$. Sei also $x \in u$ beliebig. Es gilt $x \neq \{x\}$, da $x \in u$. Dann ist aber auch $\{x\} \neq \{\{x\}\}$ nach dem Extensionalitätsaxiom und der Definition der Einermenge. Also ist $\{x\} \in u$.]

- (ii) Jedes $x \in Z_0$, $x \neq \emptyset$, hat genau ein Element.

- (iii) Z_0 ist transitiv: Ist $x \in Z_0$, so ist $x \subseteq Z_0$.

Weiter kann man eine $<$ -Relation auf Z_0 definieren durch:

$$< = \cap \{ R \subseteq Z_0 \times Z_0 \mid R \text{ ist transitiv und für alle } x \in Z_0 \text{ ist } (x, \{x\}) \in R \}.$$

Eine anspruchsvollere Übung ist:

Übung

$\langle Z_0, < \rangle$ ist eine Wohlordnung.

Hieraus gewinnt man dann Induktion und Rekursion über $\langle Z_0, < \rangle$, wobei über Z_0 rekursiv definierte Operationen $\mathcal{F} : Z_0 \rightarrow V$ im allgemeinen selbst keine Mengen mehr sind (kein Axiom garantiert uns, daß die Vereinigung von partiellen Rekursionsfunktionen „bis x “ für $x \in Z_0$ eine Menge ist, vgl. Beweis des Rekursionssatzes und die Diskussion zum Ersetzungsschema unten).

Übung

Es gilt $|Z_0| \leq |M|$ für jede Dedekind-unendliche Menge M .

Z_0 dient im Rahmen der ursprünglichen Axiome von Zermelo als Menge der natürlichen Zahlen $\mathbb{N}$. Aus der – weit nach 1908 gefundenen – Definition der Ordinalzahlen nach von Neumann und Zermelo (siehe 2.6) ergibt sich eine etwas andere Definition der natürlichen Zahlen, die man ja als Anfangsstück der Ordinalzahlen definieren möchte. Bei der Interpretation mathematischer Objekte innerhalb der Mengenlehre gibt es keine eindeutigen Lösungen. Brauchbarkeit und Natürlichkeit sind die entscheidenden Gesichtspunkte. Drängt sich eine Definition aufgrund bestimmter Überlegungen geradezu von selbst auf, so spricht man von einer „kanonischen Definition“. Und die Definition der Ordinalzahlen nach von Neumann und Zermelo ist, wie wir gesehen haben, in diesem Sinne kanonisch. Die natürlichen Zahlen sind hier:

$$0 = \emptyset, \quad 1 = \{0\}, \quad 2 = \{0, 1\}, \quad 3 = \{0, 1, 2\}, \quad 4 = \{0, 1, 2, 3\}, \dots$$

Diese Definition betont den Ordnungs- und den Mächtigkeitsaspekt der natürlichen Zahlen gleichermaßen. Jedes n hat genau n Elemente und die $\in$ -Relation ordnet die natürlichen Zahlen linear. Wählt man als Unendlichkeitsaxiom:

(UN2) *Unendlichkeitsaxiom II*

Es existiert eine Menge x , die die leere Menge als Element enthält und die mit jedem ihrer Elemente y auch $y \cup \{y\}$ als Element enthält.

– so kann man die Ordinalzahlen $0, 1, 2, \dots$ nach von Neumann und Zermelo analog definieren wie oben Z_0 definiert wurde. Man setzt:

$$\omega = \bigcap \{ U \in \mathcal{P}(x) \mid \emptyset \in U \text{ und für alle } y \in U \text{ ist } y \cup \{y\} \in U \},$$

wobei x beliebig ist wie in (UN2). Dieses ω ist dann die erste Limesordinalzahl. Man beweist, nur die Schnittdefinition verwendend, die elementaren Eigenschaften von ω , etwa: ω ist transitiv, $n \notin n$ für alle $n \in \omega$, $n \neq n+1 = n \cup \{n\}$ für alle $n \in \omega$, usw. Weiter zeigt man, daß jede nichtleere Teilmenge von ω ein kleinstes (im Sinne der $\in$ -Relation) Element hat, d.h. $\langle \omega, < \rangle = \langle \omega, \in \upharpoonright \omega \rangle$ ist eine Wohlordnung, und dies liefert Induktion und Rekursion über ω . ω ist die Menge der „natürlichen Zahlen“ innerhalb der Mengenlehre, mit der Ordinalzahldefinition nach von Neumann und Zermelo.

Was ist nun der Unterschied zwischen (UN) und (UN2)? Warum hat Zermelo Z_0 betrachtet? Es zeigt sich: (UN) und (UN2) sind äquivalent, wenn man das Ersetzungsschema zu Zermelos Axiomatik mit hinzunimmt, das wir gleich besprechen werden. In der ursprünglichen – auf den Wohlordnungssatz zugeschnittenen – Zermelo-Axiomatik kann man dagegen die Existenz von ω wie oben definiert nicht zeigen. Generell gilt, daß in einer Axiomatik ohne Ersetzungsschema nur wenige Neumann-Zermelo-Ordinalzahlen existieren, egal, welche Form man dem Unendlichkeitsaxiom gibt. Unabhängig von einem notorischen Mangel an Neumann-Zermelo-Ordinalzahlen gilt ohne Ersetzungsschema der Rekursionssatz nicht einmal für beliebige abzählbare Wohlordnungen; wir bleiben im Limeschritt hängen, und Rekursionen der Länge ω liefern i. a. schon echte Klassen. Diesem für die Mengenlehre also sehr wichtigen Axiom wenden wir uns nun zu.

Die eben aufgestellten Behauptungen über die Notwendigkeit des Ersetzungsschemas lassen sich nur durch Modellkonstruktionen beweisen, was außerhalb dieses Textes liegt. Der Leser kann aber z.B. den Beweis des Rekursionssatzes noch einmal konsultieren, nachdem er das Ersetzungsschema kennengelernt hat, um zu sehen, wo es im Beweis eingeht.

Das Ersetzungsschema

Bis auf das Auswahlaxiom, das wir am Ende der Liste besprechen, haben wir nun alle Axiome von Zermelo aus dem Jahre 1908 eingeführt. Wir kommen nun zur ersten Ergänzung des Systems von Zermelo durch das Ersetzungsschema von Abraham Fraenkel und Thoralf Skolem. Hierzu brauchen wir den Begriff einer „funktionalen Eigenschaft“ oder einer „Operation auf V “:

Ist $\mathcal{E}(x, y)$ eine Eigenschaft in zwei Variablen x und y , so heißt $\mathcal{E}(x, y)$ *funktional*, falls es für jedes x genau ein y gibt mit $\mathcal{E}(x, y)$.

$\mathcal{E}(x, y)$ ist also intuitiv eine „Funktion“ auf dem Mengenuniversum: Jedem x wird ein eindeutiges y zugeordnet, nämlich dasjenige y , für welches $\mathcal{E}(x, y)$ erfüllt ist. Beispiele sind:

$$\mathcal{E}_1(x, y) = \text{„}y \text{ ist identisch mit } x\text{“} = \text{„}y = x\text{“},$$

$$\mathcal{E}_2(x, y) = \text{„}y \text{ ist die Einermenge von } x\text{“} = \text{„}y = \{x\}\text{“},$$

$$\mathcal{E}_3(x, y) = \text{„}y \text{ ist die Potenzmenge von } x\text{“} = \text{„}y = \mathcal{P}(x)\text{“}.$$

Übung

Ist $\mathcal{E}(x, y)$ eine funktionale Eigenschaft, so existiert die Zusammenfassung $\{(x, y) \mid \mathcal{E}(x, y)\}$ nicht, d. h. es gibt keine Menge z von geordneten Paaren mit: Für alle x, y gilt: $(x, y) \in z$ gdw $\mathcal{E}(x, y)$.

[z wäre eine auf dem ganzen Mengenuniversum definierte Funktion, also keine Menge; denn andernfalls wäre nach dem Vereinigungsaxiom $\{x \mid x = x\} = \bigcup \bigcup z$ eine Menge.]

Wir können nun das Ersetzungsschema formulieren.

(ERS) Ersetzungsschema

Zu jeder funktionalen Eigenschaft $\mathcal{E}$ und jeder Menge M existiert eine Menge N , die genau diejenigen y als Elemente enthält, für welche ein $x \in M$ existiert mit $\mathcal{E}(x, y)$.

Kurz und anschaulich:

„Das Bild einer Menge unter einer funktionalen Eigenschaft ist eine Menge.“

Vorstellung ist: Wir haben eine Menge M und eine universelle Zuordnung $\mathcal{E}(x, y)$. Wir bilden nun eine Menge N , indem wir jedes Element x von M durch das eindeutige y mit $\mathcal{E}(x, y)$ „ersetzen“.

Die Anwendung des Ersetzungsschemas erzeugt intuitiv keine zu großen und damit pathologischen Objekte, da wir die Elemente einer bereits vorhandenen Menge durch andere austauschen. Die Mächtigkeit der entstehenden Menge ist also immer kleinergleich der Mächtigkeit der Ausgangsmenge.

Die Stärke des Ersetzungsschemas liegt in der Verwendung von funktionalen Eigenschaften. Ist F eine Funktion auf einer Menge M , so brauchen wir für die Existenz des Bildes $N = F''M$ von M unter F kein neues Axiom:

Übung

Sei F eine Funktion mit $M \subseteq \text{dom}(F)$. Dann existiert $N = F''M$.

[Mit Vereinigungsaxiom und Aussonderungsschema.]

Man kann das Ersetzungsschema auch als eine liberalere Form des Aussonderungsschemas betrachten:

Übung

Zeigen Sie das Aussonderungsschema mit Hilfe des Ersetzungsschemas (und der übrigen Axiome).

Wir könnten das Aussonderungsschema also aus ZFC streichen, ohne die Stärke des Axiomensystems zu vermindern. Wegen seiner Natürlichkeit und Nützlichkeit im Aufbau der Theorie hat es aber einen sicheren Platz innerhalb der Axiome.

Fraenkel (1922): „Wenn man die Cantorsche Mengenlehre, unter Ausscheidung der Antinomien und unter Verzicht auf die ihnen Raum gebende Cantorsche Mengendefinition, auf mathematisch befriedigende Grundlagen stellen will, so kommt vorläufig nur die von Herrn Zermelo gegebene Begründung in Frage. Einige das Grundgerüst dieser Begründung betreffende und z. T. es modifizierende Bemerkungen bilden den Inhalt der folgenden Zeilen ... Die überaus scharfsinnigen Untersuchungen Zermelos sollen hierdurch nicht umgestoßen, sondern nur vervollständigt und befestigt werden ...

I. Die sieben Zermeloschen Axiome reichen nicht aus zur Begründung der Mengenlehre.

Zum Nachweis dieser Behauptung diene etwa das folgende einfache Beispiel: Es sei Z_0 die [in Zermelo 1908b] definierte und als existierend nachgewiesene Menge (Zahlenreihe). Die Potenzmenge $\mathbb{I}Z_0$ (Menge aller Untermengen von Z_0) werde mit Z_1 , $\mathbb{I}Z_1$ mit Z_2 bezeichnet usw. Dann gestatten die Axiome, wie deren Durchmusterung leicht zeigt, nicht die Bildung der Menge $\{Z_0, Z_1, \dots\}$, also auch nicht die Bildung der Vereinigungsmenge. Es läßt sich daher, wenn man etwa dem Kontinuum eine Mächtigkeit $< \aleph_\omega$ zuschreibt, auf Grund der Axiome z. B. die Existenz von Mengen mit Mächtigkeit $\geq \aleph_\omega$ nicht beweisen.

Diese bisher nicht bemerkte Lücke der Zermeloschen Begründung ist durch Hinzufügung eines neuen Axioms oder Erweiterung eines vorhandenen auszufüllen ...

Ersetzungsaxiom. Ist M eine Menge und wird jedes Element von M durch ‚ein Ding des Bereiches $\mathfrak{B}$ ‘ ersetzt, so geht M wiederum in eine Menge über.

Für das oben angeführte Beispiel hat man, um die Existenz der Menge $\{Z_0, Z_1, \dots\}$ zu zeigen, auf Grund des soeben formulierten Axioms nur das Element 0 von Z_0 durch Z_0 , das Element $\{0\}$ durch Z_1 zu ersetzen usw. Man kann weiter auf die Vereinigungsmenge der so entstehenden Menge das Axiom in analoger Weise anwenden und erlangt, derart weiterschreitend, ersichtlich die erforderliche Freiheit in der Bildung von Mengen.“

Daß man tatsächlich das Ersetzungsschema nicht aus den anderen Axiomen beweisen kann, ist durch eine „Durchmusterung“ der Zermelo-Axiome nur plausibel gemacht. Der strenge Beweis dieser Behauptung fällt wie die Unabhängigkeit der Kontinuums-hypothese in den Bereich der Metamathematik: Man gibt ein Modell der Zermelo-Axiome an, in dem das Ersetzungsschema falsch ist.

Das Ersetzungsschema garantiert im Aufbau der axiomatischen Mengenlehre die Existenz vieler Neumann-Zermelo-Ordinalzahlen, kurz: Ordinalzahlen, da

die Cantor-Hausdorff-Definition hier zu ungenau ist. Weiter wird es ganz wesentlich gebraucht im Beweis des Rekursionssatzes für Wohlordnungen und allgemeiner im Beweis des Rekursionssatzes für alle Ordinalzahlen.

Wir können folglich die V_α -Hierarchie wie in 2.7 definieren. Für den Beweis, daß es für jede Menge x eine Ordinalzahl α gibt mit $x \in V_\alpha$, wird ein weiteres Axiom benötigt.

Das Fundierungsaxiom

Wir kommen nun zum Fundierungsaxiom, das kein Existenzaxiom ist, sondern dem Extensionalitätsaxiom verwandt ist: Es beschreibt die Elementrelation. Im Gegensatz zu den Existenzaxiomen sorgt es für eine Beschränkung des Mengenuniversums.

(FUN) Fundierungsaxiom

Jede nichtleere Menge x hat ein Element y , das mit x kein Element gemeinsam hat.

Also: Für alle $x \neq \emptyset$ existiert ein $y \in x$ mit der Eigenschaft $y \cap x = \emptyset$.

Das Axiom drückt aus, daß Mengen irgendwie aus „bereits vorhandenen“ anderen Mengen konstruiert werden. Das Axiom schließt Mengenzyklen aus, die die Ausgangsmenge bei einem $\in$ -Abstieg irgendwann reproduzieren:

Übung

- (i) Es gibt kein x mit $x = \{x\}$.
- (ii) Es gibt kein x mit $x \in x$.
- (iii) Es gibt keine x, y mit $x \in y \in x$.
- (iv) Es gibt keine $x_0, x_1, \dots, x_n$ mit der Eigenschaft:
 $x_0 \in x_n \in \dots \in x_2 \in x_1 \in x_0$.

Allgemeiner kann es keine unendlichen absteigenden $\in$ -Ketten geben, d. h. es gibt keine Funktion f mit $\text{dom}(f) = \mathbb{Z}_0$ (oder $= \omega$, wenn man mag) und $f(\{n\}) \in f(n)$ für alle n , d. h.

$f(\emptyset) \ni f(\{\emptyset\}) \ni f(\{\{\emptyset\}\}) \ni \dots$

Eine solche Kette kann nicht existieren, denn $x = \{f(n) \mid n \in \mathbb{Z}_0\}$ wäre dann eine Menge ohne $\in$ -minimales Element.

Das Axiom sorgt für ein scharfes Bild des Mengenuniversums V : V ist aus der leeren Menge stufenweise aufgebaut. Wie in 2.7 zeigt man, daß jedes x von einer V_α -Stufe eingefangen wird. Wird das Fundierungsaxiom außerhalb der Mengenlehre auch selten gebraucht, so ist es innerhalb der abstrakten Mengenlehre um so bedeutender. Das Universum kann entlang der Ordinalzahlen durch iterierte Anwendung der Potenzmengen- und Vereinigungsmengenoperationen vollständig durchforstet werden.

Es ist interessant, daß die beiden von Zermelo 1908 „vergessenen“ Axiome im wesentlichen nur in der Mengenlehre benötigt werden, während die Mathematik weitestgehend ohne sie auskommt. Der Grund ist, daß dort nur zwei oder dreimal überhaupt die Potenzmengenoperation verwendet wird, und nicht in unendlicher Iteration und darüber hinaus, was nur durch das Ersetzungsschema garantiert werden würde. Auch die ω -Rekursion ist im Rahmen von beschränkt vielen Potenzmengenoperationen problemlos beweisbar. Mit Fraenkels Notation: Ein Großteil der Mathematik verläuft in Z_3 oder Z_4 , und „fast alles“ etwa in Z_8 . In jedem Z_n gilt das Fundierungsaxiom, sodaß nichtfundierte Mengen gar nicht erst auftreten. (Allgemeiner gilt es in der ganzen V_α -Hierarchie.) Zermelos System erlaubt einen Beweis des Wohlordnungssatzes und des Satzes von Zermelo-Zorn, der in der Mathematik verwendet wird. Hierzu wird das Auswahlaxiom gebraucht, das wir unten besprechen werden.

Fraenkel (1922): „Hat sich [im Hinblick auf das Ersetzungsschema] der Zermelosche Mengenbegriff als zu eng für die Cantorsche Mengenlehre erwiesen, so ist er in anderer Beziehung weiter, als es die Bedürfnisse der Mathematik zu erfordern scheinen. Zunächst nämlich können unter den ‚Dingen‘ des ‚Bereiches $\mathfrak{B}$ ‘, aus denen auf Grund der Axiome die Mengen ihre Existenz herleiten, sich auch solche nichtmathematischer und überhaupt nichtbegrifflicher Herkunft befinden. Ferner läßt das Axiomensystem Raum z. B. für die von Herrn Mirimanoff als ‚ensembles extraordinaires‘ bezeichneten Mengen M von der Art, daß, wenn $M_1 \in M$ und wenn k eine beliebige natürliche Zahl bedeutet, M_k stets ein Element M_{k+1} enthält. Solche Mengen können zwar nicht auf Grund der Axiome aus den ‚unzerlegbaren‘ (d. h. keine Menge darstellenden) Dingen von $\mathfrak{B}$ aufgebaut werden, sie *können* aber in $\mathfrak{B}$ vorkommen ...

... so kann hier den angegebenen Übelständen durch ein als neuntes und letztes Axiom aufzustellendes ‚Beschränktheitsaxiom‘ abgeholfen werden, das dem Mengenbegriff oder dem Bereich $\mathfrak{B}$ den geringsten mit den übrigen verträglichen Umfang auferlegt. Verfährt man in dieser Art, so wird die Nullmenge zu dem einzigen Ding, das keine Menge [mit Elementen] ist; das genügt für alle mathematischen Zwecke und vereinfacht die Betrachtung sachlich und z. T. auch formal.“

Fraenkel plädiert also für einen Verzicht auf überflüssige Mengen. Zum einen sollen Urelemente verschwinden und das mengentheoretische Universum aus der Nullmenge aufgebaut sein. Zum anderen sollen die unendlich absteigenden $\in$ -Ketten ausgeschlossen werden. Das Ziel ist der Axiomatik einen „kategorischen Charakter“ zu geben: Die Axiome sollen den Bereich $\mathfrak{B}$, das Mengenuniversum also, wenn möglich eindeutig beschreiben. (Auf die prinzipielle Unmöglichkeit eines solchen eindeutigen Systems kommen wir im zweiten Band zu sprechen.)

Auch wenn man auf Urelemente verzichtet, kann man in der aus der leeren Menge aufgebauten Mengenlehre eine Mengenlehre mit Urelementen simulieren, indem man geeignete Mengen als Urelemente ansieht, als ein V_0^* , und über diesen „Urelementen“ eine V_α^* -Hierarchie in gewohnter Weise hochzieht. Geeignet heißt, daß während der Hierarchiebildung keine Mengen konstruiert werden, die Elemente der „Urelemente“ sind, daß also die in Wahrheit vorhandenen Elemente der „Urelemente“ dieser Hierarchie verborgen bleiben.

Von Neumann hat in seiner auf dem Funktionsbegriff aufgebauten Axiomatik ebenfalls ein Beschränkungsaxiom betrachtet.

von Neumann (1925):

„4. Es gibt kein II. Ding a [keine Funktion a] mit der folgenden Eigenschaft:
Es ist für jede endliche Ordnungszahl (d. h. ganze Zahl) n
 $[a, n + 1] \in [a, n]$ [*in unserer Schreibweise: $a(\{n\}) \in a(n)$].“*

In der Arbeit „Grenzzahlen und Mengenbereiche“ von Zermelo aus dem Jahr 1930 findet sich dann der Name „Fundierungsaxiom“, und das Axiom erscheint in der heute üblichen Formulierung ohne unendliche $\in$ -Abstiege.

Neben „Fundierungsaxiom“ ist auch die Bezeichnung „Regularitätsaxiom“ gebräuchlich.

Zermelo (1930):

„F) Axiom der Fundierung: Jede (rückschreitende) Kette von Elementen, in welcher jedes Glied Element des vorangehenden ist, bricht mit endlichem Index ab bei einem Urelement [oder bei $\emptyset$, ohne Urelemente immer bei $\emptyset$]. Oder, was gleichbedeutend ist: Jeder Teilbereich T [$\neq \emptyset$] enthält wenigstens ein Element t_0 , das kein Element t in T hat [also $t_0 \cap T = \emptyset$].

Dieses letzte Axiom, durch welches alle ‚zirkelhaften‘, namentlich auch alle ‚sich selbst enthaltenden‘, überhaupt alle ‚wurzellosen‘ Mengen ausgeschlossen werden, war bei allen praktischen Anwendungen der Mengenlehre bisher immer erfüllt, bringt also vorläufig keine wesentliche Einschränkung der Theorie.“

Einfache Modelle

Bevor wir zum Auswahlaxiom kommen, wollen wir kurz einige interessante Modelle beschreiben, die auch ohne eine genaue Definition des Modellbegriffs ein gutes Bild von der Rolle des Fundierungsaxioms und dem Verhältnis der Zermeloaxiomatik und ihrer Anreicherung durch das Ersetzungsschema liefern.

Wir können z. B. in der Axiomatik ohne Fundierungsaxiom, aber mit Ersetzung die V_α -Hierarchie definieren. Die Vereinigung aller dieser V_α bildet dann einen Bereich $\mathfrak{B}$, in dem alle Axiome gelten, einschließlich des Fundierungsaxioms: Bilden wir in $\mathfrak{B}$ erneut die V_α -Hierarchie, so ist diese Hierarchie identisch mit der alten, und schöpft $\mathfrak{B}$ voll aus. Auf diese Weise kann man zeigen: Ist die Axiomatik ohne Fundierung widerspruchsfrei, so ist sie auch mit Fundierungsaxiom widerspruchsfrei. Das einschränkende Fundierungsaxiom kann, wie erwartet, keine Widersprüche generieren.

Die Zermelo-Axiomatik (ohne Fundierung und Ersetzung) beweist, daß die übliche Zahlentheorie widerspruchsfrei ist, denn bereits in dieser abgeschwächten Axiomatik läßt sich zeigen, daß ein Modell $\langle Z_0, +, \cdot, 0, 1 \rangle$ für die Zahlentheorie existiert. Starke Theorien liefern typischerweise Modelle für schwächere

Theorien und beweisen damit die Widerspruchsfreiheit dieser Theorien. So auch für die Zermelo-Axiomatik selbst: In der Zermelo-Axiomatik erweitert um das Ersetzungsschema kann man Modelle für die Zermelo-Axiomatik angeben. Sei

$$Z^* = \bigcup_{n \in \omega} Z_n,$$

wobei wieder $Z_1 = \mathcal{P}(Z_0)$, $Z_2 = \mathcal{P}(Z_1)$, usw. Dann ist Z^* ein Modell der Zermelo-Axiomatik.

Neben Z^* haben wir in der Zermelo-Fraenkel-Axiomatik eine Fülle von Modellen für die Zermelo-Axiomatik: Für jede Limesordinalzahl $\alpha > \omega$ ist V_α ein Modell, etwa $V_{\omega+\omega}$. Man vergleiche damit, daß die volle Theorie nach dem zweiten Gödelschen Unvollständigkeitssatz kein Modell von sich selbst konstruieren kann (es sei denn sie ist widerspruchsvoll). Erst in Erweiterungen der Axiomatik um große Kardinalzahlaxiome kann man beweisen, daß viele V_α -Stufen Modelle der Zermelo-Fraenkel-Axiomatik sind, d. h. daß in vielen Limesstufen V_α nicht nur Zermelos Axiomatik, sondern auch das Ersetzungsschema gilt. Diese Stufen sind so hoch, daß ein Ersetzungsprozeß, bei dem die Elemente einer Menge $M \in V_\alpha$ durch Ordinalzahlen $\beta < \alpha$ ersetzt werden, unterhalb von α immer beschränkt bleibt. Kurz: Wir erreichen das Ende α des Ordinalzahlweges innerhalb von V_α nicht, indem wir M -viele Schritte machen, für ein beliebiges $M \in V_\alpha$.

Die Beweise dieser Behauptungen sind nicht trivial, aber auch nicht schwer. Es sind die einfachsten Beispiele der axiomatischen Mengenlehre für relative Konsistenzbeweise und Modellkonstruktionen für Teiltheorien starker Erweiterungen der Basisaxiomatik.

Damit kommen wir nun zu Zermelos Isolierung eines sehr natürlichen, und zugleich sehr starken Prinzips.

Das Auswahlaxiom

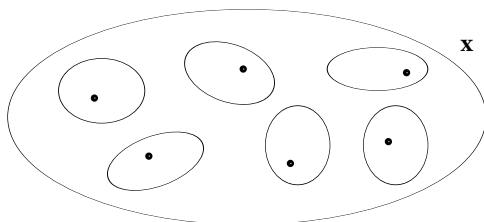
Das letzte Axiom unserer Liste, das „C“, gehört wieder der Axiomatik von Zermelo an.

(AC) Auswahlaxiom

*Ist x eine Menge, deren Elemente nichtleer und paarweise disjunkt sind,
so existiert eine Menge y ,
die mit jedem Element von x genau ein Element gemeinsam hat.*

(AC steht für engl. *axiom of choice*)

Das Auswahlaxiom ist ein Axiom im besten Sinne, da es unserer Intuition über den Mengenbegriff entspringt ...



Die Menge der Punkte bildet eine Auswahlmenge y für x

... und doch nimmt es aufgrund seines Charakters eine Sonderstellung unter den ZFC Axiomen ein. Die Auswahlmenge y , die das Axiom garantiert, ist „dunkel“, während die Mengen der anderen Existenzaxiome „hell“ sind:

Wir „sehen“ die Menge $Z_0 = \{\emptyset, \{\emptyset\}, \dots\}$, die dem Unendlichkeitsaxiom genügt. Und die anderen Existenzaxiome (LM), (PA), (VER), (AUS) und (ERS) können wir als Instanzen des Komprehensionsaxioms schreiben, und wir sehen sie vor uns, sobald wir ihre definierende Eigenschaft kennen, so etwa für $\{x, y\} = \{z \mid z = x \text{ oder } z = y\}$ oder $\mathcal{P}(x) = \{z \mid z \subseteq x\}$. (Wobei wir $\mathcal{P}(x)$ nur sehen, wenn wir einen Bereich von V überblicken können, der alle Teilmengen von x umfaßt. Das Potenzmengenaxiom ist eher sammelnd denn erzeugend, und in diesem Sinne ist die Potenzmenge einer Menge halbhell und halbdunkel.)

Ein vergleichbares Sehen einer Auswahlmenge y für x ist dagegen nur in Spezialfällen möglich. Wir können die Menge y im allgemeinen nicht in der Form $y = \{z \mid \mathcal{E}(z)\}$ schreiben, es sei denn, alle $a \in x$ besitzen ausgezeichnete Elemente: Gibt es eine Eigenschaft $\mathcal{E}(z)$ mit:

„für alle $a \in x$ existiert genau ein $z \in a$ mit $\mathcal{E}(z)$ “,

so können wir eine Auswahlmenge y mit Hilfe von $\mathcal{E}(z)$ sichtbar machen:

$$y = \{z \in \bigcup x \mid \mathcal{E}(z)\}.$$

y existiert aufgrund des Vereinigungsaxioms und des Aussonderungsschemas.

Unsere Intuition gibt uns keinen Hinweis, daß stets ausgezeichnete Elemente innerhalb jedes $a \in x$ vorhanden sind. Und somit brauchen wir ein neues Axiom, um eine Auswahlmenge für beliebige x garantieren zu können. Die Auswahlmenge, die uns das Axiom liefert, ist allerdings in keiner Weise mehr eindeutig bestimmt, sie ist zufällig, sie ist „dunkel“. In einem sehr natürlichen Modell der Mengenlehre, dem konstruktiblen Universum L von Gödel, sind tatsächlich immer ausgezeichnete Elemente vorhanden, und wir können das Auswahlaxiom dann in diesem Modell auf der Basis der übrigen Axiome beweisen, da wir in diesem Modell für jedes x eine sichtbare Auswahlmenge definieren können.

Berühmt ist die populäre Formulierung durch Bertrand Russell: Gegeben eine unendliche Menge von Schuhpaaren, können wir ohne Auswahlaxiom eine Menge definieren, die von jedem Paar genau einen Schuh enthält, z. B. die Menge aller linken Schuhe. Dagegen brauchen wir für eine unendliche Menge von Sockenpaaren das Auswahlaxiom, da wir keine Möglichkeit haben, zwischen linken und rechten Socken zu unterscheiden.

Zermelo hat sein Axiom der „simultanen Auswahl“ in seinen Beweisen des Wohlordnungssatzes von 1904 und 1908 substantiell verwendet, und es als allgemeines Prinzip isoliert. Sein Beweis von 1904 hat eine kontroverse und weitgehend irrationale Diskussion um die Legitimation einer „simultanen Auswahl“ hervorgerufen. Heute ist das Auswahlaxiom ein fundamentaler und unumstrittener Bestandteil der Axiomatik der Mengenlehre. Es gibt aber mittlerweile interessante Modelle der Mengenlehre, in denen lediglich abgeschwächte Formen des Auswahlaxioms, sonst aber alle übrigen Axiome – und andere – gelten.

Die Formulierung von (AC) über paarweise disjunkte Mengen ist sehr einfach und anschaulich. Häufig werden aber Auswahlfunktionen gebraucht:

Übung

Zeigen Sie, daß (AC) auf der Basis der übrigen Axiome jeweils äquivalent ist zu:

- (i) Ist M eine Menge mit $\emptyset \notin M$, so existiert eine Funktion $f: M \rightarrow \bigcup M$ mit $f(x) \in x$ für alle $x \in M$.
- (ii) Ist g eine Funktion mit $g(x) \neq \emptyset$ für alle $x \in \text{dom}(g)$, so existiert eine Funktion f mit $\text{dom}(f) = \text{dom}(g)$ und $f(x) \in g(x)$ für alle $x \in \text{dom}(f)$.
- (iii) Jede Äquivalenzrelation besitzt ein vollständiges Repräsentantensystem.

Die Aussage (ii) kann man auch so formulieren: Es gilt $\bigcap_{i \in I} A_i \neq \emptyset$ für alle Mengen $I \neq \emptyset$, falls $A_i \neq \emptyset$ für alle $i \in I$.

Wissen wir, daß eine Menge M nichtleer ist, so ist für „Sei also $x \in M$ beliebig.“ natürlich kein Auswahlaxiom notwendig. Allgemein zeigt man durch Induktion nach $|M|$ ohne Auswahlaxiom, daß für $\mathbb{N}$ -endliche Mengen M mit $\emptyset \notin M$ immer eine Auswahlfunktion wie in (i) existiert, oder daß das kartesische Produkt von endlich vielen nichtleeren Mengen immer nichtleer ist.

Das Auswahlaxiom wird für die meisten weitergehenden Resultate über Mächtigkeiten gebraucht. Eine Ausnahme bildet der Satz von Cantor-Bernstein, der sich ohne Verwendung von (AC) beweisen läßt. Dagegen ist für einen Beweis des Vergleichbarkeitssatzes das Auswahlaxiom unverzichtbar. Selbst den Satz, daß die abzählbare Vereinigung von abzählbaren Mengen wieder abzählbar ist, kann man nicht ohne eine Form von (AC) beweisen. In den beiden ersten Abschnitten haben wir das Auswahlaxiom darüber hinaus an mehreren Stellen wesentlich benutzt, und durch „ein ...“ darauf hingewiesen. Diese Ausdrucksweise „ $g(y) = \text{ein } x \in M \text{ mit } \mathcal{E}(x, y)$ “ ist in der axiomatischen Mengenlehre nur eine bequeme Sprechweise für $g(y) = f(\{x \in M \mid \mathcal{E}(x, y)\})$, wobei f eine Auswahlfunktion wie in (i) der Übung oben ist, die man zu Beginn des Beweises auf $\mathcal{P}(M) - \{\emptyset\}$ fixiert (vgl. etwa Zermelos Beweis des Wohlordnungssatzes). Dieses Präludium der Fixierung einer Auswahlfunktion lenkt aber i. a. nur ab, es genügt, wenn man während des Beweises sicherstellt, daß für alle y ein $x \in M$ mit $\mathcal{E}(x, y)$ existiert. Man führt also Beweise üblicherweise ganz wie gehabt mit „ein ...“.

Vorsicht ist geboten bei der Rekursion über alle Ordinalzahlen. Dort ist $\mathcal{G}(\alpha) = \text{„ein ...“}$ i.a. nicht erlaubt, da wir in ZFC i.a. keine definierbare Auswahloperation $\mathcal{F}$ auf $V - \{\emptyset\}$ haben. Im Beweis des Rekursionssatzes wird eine feste Operation $\mathcal{F}$ gebraucht. Läuft die Rekursion über eine Wohlordnung (d. h. eine Menge), so ist „ein ...“ möglich, da man dann eine Auswahlfunktion benutzen kann. Allgemein macht z. B. $\mathcal{G}(\alpha) = \text{„ein } x \in V_{\alpha+1} - V_{\alpha}\text{“}$ in ZFC keinen Sinn.

Übung

Lokalisieren Sie die Verwendung von (AC) in den Beweisen von:

- (i) $|M| \leq |N|$ oder $|N| \leq |M|$ für alle Mengen M, N .
- (ii) $\bigcup \{A_n \mid n \in \mathbb{N}\}$ ist abzählbar, falls alle A_n abzählbar sind.
- (iii) „ $|M| \leq |N|$ “ *gdw* „es existiert ein $f: N \rightarrow M$ surjektiv“.
- (iv) M ist (Dedekind-)endlich *gdw* „ $|M| = |\bar{n}|$ für ein $n \in \mathbb{N}$ “.

Das Auswahlaxiom kommt in vielen Bereichen der Mathematik wesentlich zum Einsatz, zumeist in Gestalt eines Maximalprinzips.

Hier noch einmal eine kleine Liste für den notwendigen Einsatz des Auswahlaxioms, vgl. auch die Diskussion von Maximalprinzipien in 2.5: „jede Äquivalenzrelation hat ein vollständiges Repräsentantensystem“, „jeder Vektorraum hat eine Basis“, Satz von Hahn-Banach, Satz von Tychonov, Existenz- und Eindeutigkeitssatz des algebraischen Abschlusses eines Körpers, Kompaktheitssatz der Logik, Satz von Tarski über die Existenz von Ultrafiltern, Primidealtheorem für Boolesche Algebren.

Es ist zudem verantwortlich für die Existenz von nicht Lebesgue-meßbaren Teilmengen von $\mathbb{R}$, und für andere „paradoxe“ Konstruktionen, die der Maßtheorie i. a. den Zugang zur vollen Potenzmenge einer Menge, deren Teilmengen sie gerne messen möchte, verwehren.

In der Mengenlehre ist (AC) unter anderem äquivalent (über den anderen Axiomen) zum Wohlordnungssatz, zum Vergleichbarkeitssatz, zum Satz von Zermelo-Zorn und zum Multiplikationssatz.

Zermelo (1908b):

„**Axiom VI.** Ist T eine Menge, deren sämtliche Elemente von 0 verschiedene Mengen und untereinander elementenfremd sind, so enthält ihre Vereinigung $\bigcup T$ mindestens eine Untermenge S_1 , welche mit jedem Element von T ein und nur ein Element gemein hat.

(Axiom der Auswahl.)“

Damit sind alle Axiome von ZFC vorgestellt. Es gibt alternative Axiomatisierungen – das System NBG von Neumann-Bernays-Gödel und das System MK von Morse-Kelley –, die sich aber von ZFC nicht durch den Gehalt der Axiome unterscheiden, sondern durch die liberalere Behandlung von Klassen, also von Zusammenfassungen $\{x \mid \mathcal{E}(x)\}$. In diesen Systemen ist jedes Objekt eine Klasse, und die Klassen, die Elemente einer anderen Klasse sind, sind die Mengen. Jede Menge ist eine Klasse, aber nicht umgekehrt. Die Mengen sind dann gerade die

„kleinen“ Klassen des Systems. In diesen Systemen kann man also offiziell über Objekte $\{x \mid \mathcal{E}(x)\}$ reden, während wir Klassen in ZFC nur als bequeme Sprechweise verwenden werden (hierzu und zu NBG und MK siehe 3.3). Insgesamt verfolgen diese alternativen Systeme eine zur Zermelo-Fraenkel-Axiomatik recht ähnliche Interpretation der Paradoxien.

Wir listen die Axiome von ZFC mit einer Beschreibung ihres Charakters noch einmal in anderer Anordnung auf.

Die ZFC-Axiome

(EXT)	Extensionalitätsaxiom	Beschreibung von $=, \in$
(FUN)	Fundierungsaxiom	Beschreibung von $\in$
(LM)	Existenz der leeren Menge	elementares Existenzaxiom
(PA)	Paarmengenaxiom	elementares Existenzaxiom
(VER)	Vereinigungsmengenaxiom	elementares Existenzaxiom
(AUS)	Aussonderungsschema	elementares Existenzaxiom
(UN)	Unendlichkeitsaxiom	starkes Existenzaxiom
(ERS)	Ersetzungsschema	starkes Existenzaxiom
(POT)	Potenzmengenaxiom	starkes Existenzaxiom
(AC)	Auswahlaxiom	starkes Existenzaxiom

Gängige Bezeichnungen für bestimmte Teilsysteme von ZFC sind:

Z = ZFC ohne (ERS), (FUN), (AC),

ZF = ZFC ohne (AC),

ZFC^- = ZFC ohne (POT),

ZF^- = ZF ohne (POT).

ZFC ist das Ergebnis einer natürlichen und sorgfältigen Analyse der intuitiven Begriffe Menge und Element. Diese Analyse ist dabei an den Bedürfnissen der

mathematischen Praxis orientiert – Zermelo ließ sich bei der Aufstellung seiner Axiome von seinem Beweis des Wohlordnungssatzes leiten – und ist frei von philosophischen Dogmen und Verboten. Alle wesentlichen Ergebnisse der Cantorschen Mengenlehre bleiben erhalten und die Antinomien des vollen Komprehensionsprinzips lösen sich auf. Die auf der Elementrelation basierende Sprache erweist sich zudem bei all ihrer Einfachheit als suggestiv und ausdrucksstark, und trägt wesentlich zum Erfolg der axiomatischen Mengenlehre bei.

Es hat sich gezeigt, daß diese Analyse auch vollständig unsere Intuition erschöpft: ZFC reicht für die ganze Mathematik als Basistheorie aus, und dies spricht dafür, daß keine einfachen und unmittelbar intuitiven Axiome mehr zu ZFC hinzukommen werden. ZFC scheint auch korrekt zu sein: Bis zum heutigen Tag sind in ZFC keine Widersprüche festgestellt worden. Ein mathematischer Beweis der Widerspruchsfreiheit von ZFC läßt sich, wie erwähnt, aufgrund des zweiten Gödelschen Unvollständigkeitssatzes generell nicht erbringen (es sei denn, ZFC ist widerspruchsvoll)

Denkbar ist allenfalls, daß das Konzept eines fertigen unendlichen Objekts für sich schon zu Widersprüchen führt. Das mittlerweile sehr umfangreiche Gebäude der Mengenlehre zeigt aber, daß ein solcher Widerspruch, wenn es ihn überhaupt gibt, wahrscheinlich sehr tief liegt, im Gegensatz etwa zu den klassischen Paradoxien. Wir müssen mit der Möglichkeit eines Widerspruchs in ZFC leben.

Cantor hat Ende des 19. Jahrhunderts aus seiner Mengendefinition Prinzipien über die Existenz von Mengen abgeleitet. Er denkt zu dieser Zeit intensiv über seine Unterscheidung beliebiger Vielheiten (Klassen) in „fertige Mengen/konsistente Vielheiten“ (Mengen) und „inkonsistente Vielheiten/absolut unendliche Vielheiten“ (echte Klassen) nach. In einem Brief an Hilbert im Jahre 1898 isoliert er einige Prinzipien, die beschreiben, wann eine Vielheit als fertig und damit als Menge gelten darf. Eine klare Antwort auf die Vielheits-Frage konnte er nicht geben. Der folgende Auszug aus dem Brief an Hilbert zeigt aber einmal mehr sein tiefes Verständnis des Mengenbegriffs: In Folgerung II dieses Briefes formuliert er bereits das Ersetzungsaxiom, das erst Jahrzehnte später zur Axiomatik von Zermelo als „vergessenes Axiom“ hinzukam, und das beim axiomatischen Aufbau der Mengenlehre heute eine wichtige Rolle spielt. In Folgerung IV formuliert er das Potenzmengenaxiom. Interessant ist, daß ihm die Begründung von IV mit Hilfe seiner Mengendefinition bereits zwei Tage später nicht mehr unproblematisch erscheint.

Georg Cantor an David Hilbert über „fertige Mengen“

„Lieber Herr Kollege,

Zu dem, was ich Ihnen am 6^{ten} Oktober geschrieben, möchte ich noch einiges hinzufügen, mit dem Ersuchen, Alles was ich Ihnen schreibe, Ihrer Kritik zu unterwerfen.

Aus der Definition:

„Unter einer *fertigen Menge* verstehe man jede Vielheit, bei welcher alle Elemente *ohne Widerspruch* als *zusammenseiend* und daher als *ein Ding für sich* gedacht werden können.“

ergeben sich mancherlei Sätze, unter Anderm diese:

- I „Ist M eine fertige Menge, so ist auch jede Teilmenge von M eine fertige Menge.“
- II „Substituiert man in einer fertigen Menge an Stelle der Elemente fertige Mengen, so ist die hieraus resultierende Vielheit eine fertige Menge.“
- III „Ist von zwei äquivalenten [gleichmächtigen] Vielheiten die eine eine fertige Menge, so ist es auch die andere.“
- IV „Die Vielheit *aller Teilmengen* einer fertigen Menge M ist eine fertige Menge.“
Denn alle Teilmengen von M sind ‚zusammen‘ in M enthalten; der Umstand, daß sie sich teilweise decken, schadet hieran nichts.“

(Georg Cantor an David Hilbert, Brief vom 10.10.1898. In: Cantor, 1991, Briefe)

„Lieber Herr Kollege,

Unter Bezugnahme auf mein Schreiben vom 10^{ten}, stellt sich bei genauerer Erwägung heraus, daß der Beweis des Satzes IV keineswegs so leicht geht. Der Umstand, daß die *Elemente* der ‚Vielheit *aller Teilmengen* einer fertigen Menge‘ sich teilweise decken, macht ihn illusorisch. In die Definition der fertigen Menge wird die Voraussetzung des *Getrenntseins* resp. *Unabhängigseins* der Elemente als *wesentlich* aufzunehmen sein.

Hoffentlich führt unsere Diskussion zur allmählichen Klärung der Schwierigkeiten.

Mit bestem Gruß

Ihr

G.C.“

(Georg Cantor an David Hilbert, Brief vom 12.10.1898. In: Cantor, 1991, Briefe)

2. Die Sprache der Mengenlehre

In diesem Kapitel stellen wir den formalen Rahmen für die axiomatische Mengenlehre auf. Wir konstruieren eine mathematische Kunstsprache, in der wir die Axiome der Mengenlehre und allgemeiner beliebige mengentheoretische Aussagen formulieren können.

Die formale Welt ist die Welt größtmöglicher mathematischer Präzision. Da die in ihr herrschende Form der Genauigkeit für die übliche menschliche Mathematik unfruchtbar ist – formale Beweise führt nur eine Maschine, und das bislang recht dürftig –, liegt diese Welt normalerweise zurecht unter der Oberfläche. Dieser mathematische Keller hat aber bei all seiner Unwirtlichkeit doch einige interessante Aspekte zu bieten. Für metamathematische Resultate ist ein formales System unverzichtbar, auch wenn die Beweise dieser Resultate dann schließlich wieder in der flexibleren mathematischen Umgangssprache geführt werden.

Beispiel einer Analyse einer Eigenschaft

Wir hatten oftmals Ausdrücke verwendet der Form: „Sei $\mathcal{E}(x)$ eine Eigenschaft“, etwa im Aussonderungs- und Ersetzungsaxiom. Was genau ist nun eine solche Eigenschaft? Zermelo spricht in seiner Axiomatik im Aussonderungsschema statt von Eigenschaften von „definiten Aussagen“. Die Kritik an seinem unscharfen Eigenschaftsbegriff blieb nicht aus. Der von Thoralf Skolem in folgendem Auszug aus einer Rede von 1922 vorgeschlagene Weg führt direkt zu der heute üblichen Präzisierung der Begriffe einer mathematischen Eigenschaft und einer mathematischen Aussage.

Skolem (1922): „Bisher hat, soweit mir bekannt nur ein solches Axiomensystem eine ziemlich allgemeine Anerkennung gefunden, nämlich das von E. ZERMELO aufgestellte ...

ZERMELO betrachtet einen Bereich B von Dingen, unter denen die Mengen einen Teil bilden. Zwischen diesen Dingen bestehen Beziehungen der Form $a \in b$ (a Element von b) und $a = b$. Für diesen Bereich sollen dann 7 Axiome erfüllt sein ...

Ein sehr unvollkommener Punkt bei ZERMELO ist der Begriff ‚definite Aussage‘. Die Erklärungen ZERMELOS darüber wird wohl keiner befriedigend finden. Soweit mir bekannt, hat niemand versucht, diesen Begriff streng zu formulieren, was sehr sonderbar ist, da dies leicht zu machen ist und zwar in einer sehr natürlichen Weise, die sich von selbst ergibt. Um dies zu erklären – und auch mit Rücksicht auf die späteren Betrachtungen – erwähne ich hier die 5 Grundoperationen der mathematischen Logik, wobei ich die Bezeichnungen E. SCHRÖDERS (Algebra der Logik) benutze:

1 $\times$. Die Konjunktion. Durch einen Punkt oder Nebeneinanderstellung bezeichnet.

1. Die Disjunktion. Durch das Zeichen $+$ bezeichnet.
 2. Die Negation. Durch einen Strich bezeichnet, der oberhalb des zu negierenden Ausdrucks geschrieben wird.
 - 3_×. In jedem Falle Gültigkeit. Durch das Zeichen Π bezeichnet.
 - 3₊. In mindestens einem Falle Gültigkeit. Durch das Zeichen Σ bezeichnet.
- Bekanntlich braucht man eigentlich nur 3 dieser 5 Operationen, weil $1_{\times}$ und 1_{+} ebenso wie $3_{\times}$ und 3_{+} mit Hilfe von 2 auseinander ableitbar sind.

Unter einer definiten Aussage kann man jetzt einen endlichen Ausdruck verstehen, der von Elementaraussagen der Form $a \in b$ oder $a = b$ mit Hilfe der 5 genannten Operationen [Konjunktion, Disjunktion, Negation, Allquantifizierung, Existenzquantifizierung] aufgebaut ist.

Das ist ein vollkommen klarer Begriff und hinreichend umfassend, um alle gewöhnlichen mengentheoretischen Beweise durchführen zu können. Ich lege deshalb diese Auffassung hier zu Grunde.“

Die Ausdrücke „ $a \in b$ “ und „ $a = b$ “ werden also durch logische Kombination und Quantifizierung zu mengentheoretischen Aussagen. Man verwendet heute andere Zeichen als Schröder und Skolem – ein Querstrich über einer ganzen Formel als Negation ist völlig untragbar –, aber was Skolem hier beschreibt und vorschlägt, ist die Erziehung der Mengenlehre in der Sprache der Prädikatenlogik, ganz so, wie es heute der Brauch ist.

Bevor wir die Syntax dieses Esperantos der Mengenlehre genauer einführen, analysieren wir eine spezielle Eigenschaft $\mathcal{E}(R)$, nämlich:

- (1) *R ist eine (zweistellige) Relation.*

Wir schreiben „ R ist eine Relation“ ausführlicher als:

- (2) *Für alle $x \in R$ gilt: x ist ein geordnetes Paar.*

Dies wiederum können wir in mehreren Schritten immer ausführlicher schreiben:

- (3) *Für alle $x \in R$ gibt es y, z mit $x = \{\{y\}, \{y, z\}\}$.*

- (4) *Für alle $x \in R$ gibt es y, z mit:*

$\{y\} \in x$ und $\{y, z\} \in x$

und:

für alle $z' \in x$ gilt: $z' = \{y\}$ oder $z' = \{y, z\}$.

- (5) *Für alle $x \in R$ gibt es y, z mit:*

es gibt ein $a \in x$ mit: für alle u gilt: $u \in a$ gdw $u = y$

und:

es gibt ein $b \in x$ mit: für alle u gilt: $u \in b$ gdw $u = y$ oder $u = z$

und:

für alle $z' \in x$ gilt:

für alle u gilt: $u \in z'$ gdw $u = y$

oder

für alle u gilt: $u \in z'$ gdw $u = y$ oder $u = z$.

In dieser Form haben wir zwar die Les- und Verstehbarkeit von
„R ist eine Relation“

fast verloren, jedoch haben wir diese Aussage nun in einen gleichwertigen Ausdruck verwandelt, der lediglich enthält:

- (i) *Variablen* $x, y, z, z', R, \dots$
- (ii) *die logischen Quantoren* „für alle ... gilt“ und „es gibt ... mit“,
- (iii) *die logischen Verknüpfungen* „gdw“, „oder“, „und“,
- (iv) *atomare Ausdrücke* der Form „ $x = y$ “, „ $x \in y$ “.

Im obigen Beispiel kommen alle Variablen x mit Ausnahme von R in der Form „für alle x “ oder „es gibt ein x “ vor: Sie sind *gebunden*, während R *frei* ist.

Mit (i) – (iv) haben wir schon die wesentlichen Elemente unserer Sprache beisammen. Wir werden Eigenschaften $\mathcal{E}$ und mathematische Aussagen selten derart elementar ausschreiben, aber wir könnten es im Prinzip tun.

Wir können auch Exaktheit *und* Lesbarkeit zugleich erreichen, wenn wir bereits definierte Begriffe zulassen. Wenn wir z. B. jetzt „ f ist eine Funktion“ genau ausschreiben wollen, so schreiben wir „ f ist eine Relation *und* f ist rechts-eindeutig“, und müssen uns nur noch um den zweiten Teil kümmern, da wir den Begriff „ f ist eine Relation“ schon exakt eingeführt haben.

Ein subtiler Punkt ist hier die Variablenkollision: Oben haben wir „ R ist eine Relation“ definiert. Wenn wir nun „ f ist eine Relation“ ausschreiben wollen, müssen wir lediglich in (5) die Variable R überall durch die Variable f ersetzen. Wenn wir nun aber „ x ist eine Relation“ exakt definieren wollen, kommt es in (5) zu einer Kollision der Variablen, da x dort als Hilfsvariable verwendet wird. Man kann derartige Kollisionen immer umgehen, indem man zunächst alle Hilfsvariablen geeignet umbenennt, und dann überall die Substitution durchführt.

Auch in obigem Beispiel wäre das Verfahren des Rückgriffs auf bereits Definiertes nützlich gewesen. Zuerst schreiben wir ausführlich auf, was ein geordnetes Paar ist (und hierfür noch früher, was $\{y\}$ und $\{y, z\}$ ist). Nach diesen Vorbereitungen ist „ f ist Relation“ durch „für alle $x \in f$ gilt: x ist ein geordnetes Paar“ exakt definiert. Man kann wie erwartet zeigen, daß solche Anreicherungen der Sprache um definierte Begriffe keine neue Ausdruckstärke mit sich bringen, da sich die eingeführten Begriffe wieder eliminieren lassen – wir können sie einfach als Abkürzungen ansehen.

Wir erhalten so kumulativ einen Bestand an exakt definierten mathematischen Begriffen, mit dem wir schließlich fast so frei umgehen können wie zuvor, nur daß wir jetzt Definitionen zur Verfügung haben, die sich letztendlich auf „ $x \in y$ “ und „ $x = y$ “ zurückführen lassen. Vor allem aber erhalten wir eine Definition dessen, was überhaupt eine Eigenschaft von x ist: Eine Eigenschaft von x ist einfach ein syntaktisch korrekter Ausdruck unserer Kunstsprache, der lediglich x als „freie Variable“ enthält, d. h. über alle anderen in dem Ausdruck vorkommenden Variablen y, z , usw. wird in der Form „für alle y “ oder „es gibt ein y “, usw. quantifiziert. Beliebige syntaktisch korrekte Ausdrücke heißen *Formeln*. Weiter ist eine *Aussage* oder ein *Satz* eine Formel unserer Kunstsprache ohne freie Variable.

Übung

Schreiben Sie die Eigenschaften

- (i) „ $f : A \rightarrow B$ ist eine Bijektion“,
- (ii) „ x ist unendlich“ (Dedekind-Definition),
- (iii) „ x ist abzählbar“,
- (iv) „ y ist die Potenzmenge von x “

wie im Beispiel für „ R ist Relation“ aus, unter Verwendung von bereits definierten Eigenschaften. ((i) ist ein Ausdruck mit drei freien Variablen: f, A, B ; (iv) hat die freien Variablen x, y .) Ebenso für die Aussagen:

- (v) „für alle x existiert $\mathcal{P}(x)$ “,
 - (vi) „für alle x existiert keine Bijektion $f : x \rightarrow \mathcal{P}(x)$ “.
- [„nicht“ gilt wie „und“, „oder“, ... als Bestandteil der Sprache.]

Metaebene und Metamathematik

Bei der Festsetzung und Analyse des formalen Rahmens für die Mengenlehre oder allgemeiner für die Mathematik befinden wir uns auf einer anderen Ebene als üblich. Wir diskutieren keine inhaltlichen Fragen über Mengen oder andere mathematischen Objekte, sondern formulieren die Sprache und die Regeln einer solchen Diskussion. Diese zweite Ebene wird als „Metaebene“ bezeichnet. Die Metaebene beschäftigt sich – in mathematischer Weise – mit der Mathematik selbst.

In den ersten beiden Abschnitten wurde ein Rahmen für die Mathematik überhaupt nicht weiter diskutiert – die Mengenlehre haben wir nach dem Schema Definition, Satz, Beweis entwickelt in einer für die Mathematik geeigneten Form der Umgangssprache, in der z. B. „es gibt ein x “ per Konvention bedeutet „es gibt mindestens ein x “. Was ein Satz, eine Definition, ein Beweis ist, wurde nie besprochen – man lernt diese Dinge durch Nachahmung.

Nun wollen wir die Mengenlehre als mathematische Theorie betrachten, und dabei z. B. die Frage beantworten, was eine mengentheoretische Eigenschaft ist. Unser Vorgehen ist dabei dem der üblichen Mathematik sehr ähnlich. Im Unterschied zur üblichen Mathematik besitzt die Metaebene jedoch einen gänzlich finiten Charakter, ihre Gegenstände sind konkrete Objekte, nämlich Zeichenreihen und Listen von Zeichenreihen. (Oder in einer akustischen Kultur: Lautfolgen und Sequenzen von Lautfolgen.)

Auf der Metaebene haben wir ein gewisses Maß an mathematischen Hilfsmitteln zur Verfügung, etwa die (metamathematischen) natürlichen Zahlen, einfache Arithmetik mit diesen Zahlen, usw. Wir verzichten hier darauf, genau aufzulisten, welche mathematischen Hilfsmittel wir auf der Metaebene verwenden. Es genügt uns hier, darauf zu achten, daß wir den finiten Charakter der Metaebene an keiner Stelle sprengen. Fertige unendliche Objekte werden an keiner Stelle benötigt. Auch Aussagen wie „für alle metamathematischen n gilt ...“ werden nie wirklich

für alle n benötigt – wie auch in einer endlichen Welt ? –, sondern sagen uns, was für ein beliebiges konkretes n gilt, und die zugehörigen Beweise zeigen uns, warum dies für jedes n so sein muß. Wir lassen Induktion und Rekursion als eine Sprechweise zu, die letztendlich lediglich Schemata generiert. Alles auf der Metaebene ist effektiv und in konkreten Fällen auflösbar. Der Leser betrachte etwa: „Jeder Satz, der aus den Worten ‚Aal‘, ‚Waage‘ und ‚rabenschwarz‘ zusammengesetzt ist, hat eine gerade Anzahl von Buchstaben ‚a‘.“ Um dies einzusehen, brauchen wir nicht wirklich eine Metatheorie, die über unendliche Mengen von Sätzen redet, Funktionen auf den metamathematischen natürlichen Zahlen studiert, usw. Eine Induktion über die Anzahl n der Wörter eines Satzes bestehend aus den drei Wörtern zeigt die Behauptung, und ist völlig gleichwertig zu: „Streiche das letzte Wort und addiere 2 zur bisherigen Anzahl der gefundenen a's. Wiederhole das Verfahren, bis kein Wort mehr übrig ist.“

Auch von endlichen oder effektiv unendlichen Mengen auf der Metaebene zu reden ist problemlos. Um die Ebenen auseinanderzuhalten, und um den effektiven Charakter der Metaebene zu betonen, reden wir aber bevorzugt von Listen.

Die Metamathematik besteht aus den Ergebnissen, die sich über eine formalisierte Mathematik gewinnen lassen. Beispiele für Resultate der Metamathematik sind:

„Die Kontinuumshypothese ist in ZFC weder beweisbar noch widerlegbar, wenn ZFC widerspruchsfrei ist.“

„Die Widerspruchsfreiheit von ZFC kann innerhalb von ZFC nicht bewiesen werden, wenn ZFC widerspruchsfrei ist.“

„Wenn ZF widerspruchsfrei ist, so ist auch ZFC widerspruchsfrei.“

Durch die im ersten Abschnitt geschilderte Verwendung von Modellen wird es möglich, daß z. B. der Beweis der ersten dieser drei metamathematischen Aussagen in einer Weise geführt wird, die sich von üblichen mathematischen Beweisen kaum unterscheidet. Keineswegs wird das Ziel durch Analyse der Syntax erreicht, sondern durch Konstruktion von zwei Modellen von ZFC, in denen (CH) einmal wahr und einmal falsch ist. Aus der Korrektheit des Modellbegriffs für den formalen Rahmen folgt dann die Unmöglichkeit eines formalen Beweises von (CH) oder non (CH). Auf diese Weise kann die menschliche Seite der Mathematik in die Metaebene integriert werden. Letztendlich ist aber jeder derartige Beweis, der eine metamathematische Aussage durch eine Flucht in die Objektwelt zeigt, so effektiv wie das Zählen des Buchstabens a in einem beliebigen Satz. Er liefert im Fall der Nichtbeweisbarkeit und Nichtwiderlegbarkeit der Kontinuumshypothese ein komplexes, aber dennoch konkretes Verfahren, das folgendes leistet: Es nimmt einen Beweis in ZFC entgegen. Ist dieser Beweis ein Beweis von (CH) oder von non (CH), so wird dieser Beweis in einen Beweis von „ $0 = 1$ “ in ZFC umgewandelt.

Die Sprache $\mathcal{L}$ der Mengenlehre

Wir werden nun die Kunstsprache der Mengenlehre definieren. Syntaktische Angelegenheiten sind immer ein wenig trocken, und zugleich gehen sie mit Definitionsfluten einher. Wir bemühen uns um Kürze, und verzichten bewußt auf eine allzu penible Darstellung.

Es ist üblich, auch auf der Metaebene die Begriffe *Definition*, *Satz*, *Beweis* zu verwenden, obwohl sie vielleicht eher *Festsetzung*, *Faktum* und *Nachweis* genannt werden sollten, da es um Aussagen über die Realität geht. Man vergleiche den Unterschied von „der Ausdruck $\{x \exists y\}$ enthält den Buchstaben z nicht“, ein Faktum, und „5 ist eine Primzahl“, ein mathematischer Satz.

Die Zeichen von $\mathcal{L}$

Definition (Zeichen)

Die Zeichen von $\mathcal{L}$ bestehen aus:

- (i) Variablenzeichen $v_0, v_1, v_2, \dots$
- (ii) den Relationszeichen $=$ und $\in$ [gelesen: „gleich“ und „Element von“],
- (iii) den Quantorenzeichen $\forall$ und $\exists$ [gelesen: „für alle“ und „es gibt“],
- (iv) den Junktorenzeichen $\neg, \wedge, \vee, \rightarrow, \leftrightarrow$ [non, und, oder, folgt, genau-dann-wenn],
- (v) Klammersymbolen $($ und $)$.

In (i) meinen wir, daß wir einen beliebig großen Vorrat an Variablenzeichen zur Verfügung haben. Wir indizieren diese Zeichen mit (metamathematischen) natürlichen Zahlen.

Die Ausdrücke und Formeln von $\mathcal{L}$

Definition (Ausdrücke)

Ein Ausdruck von $\mathcal{L}$ ist eine Liste $s_0 s_1 s_2 \dots s_n$, in der jedes s_i ein Zeichen von $\mathcal{L}$ ist.

So ist z. B. $(v_1 = \leftrightarrow))\forall$ ein Ausdruck von $\mathcal{L}$.

Wir verwenden auch alle Buchstaben, die wir bislang für Mengen verwendet haben, als Variablenzeichen. Dies dient lediglich der Bequemlichkeit. Schreiben wir im folgenden z. B. „sei A ein Ausdruck der Form $x = y$ “, so sind mit x, y immer Variablenzeichen gemeint, d. h. es gibt i, j , sodaß A die Zeichenkette $v_i = v_j$ ist.

Definition (*Primformeln und Formeln*)

Die *Primformeln* oder *atomaren Formeln* von $\mathcal{L}$ sind festgelegt durch:

- (i) Jeder Ausdruck der Form „ $x = y$ “ oder „ $x \in y$ “ ist eine Primformel.
- (ii) Keine weiteren Ausdrücke sind Primformeln.

Die *Formeln* von $\mathcal{L}$ sind festgelegt durch:

- (i) Jede Primformel ist eine Formel.
- (ii) Ist φ eine Formel, so ist auch $\neg\varphi$ eine Formel.
- (iii) Sind φ und ψ Formeln, so sind auch $(\varphi \wedge \psi)$, $(\varphi \vee \psi)$, $(\varphi \rightarrow \psi)$ und $(\varphi \leftrightarrow \psi)$ Formeln.
- (iv) Ist x eine Variable und φ eine Formel, so sind auch $\forall x\varphi$ und $\exists x\varphi$ Formeln.
- (v) Keine weiteren Ausdrücke sind Formeln.

Im folgenden verwenden wir zumeist kleine griechische Buchstaben für Formeln.

Damit sind die wesentlichen Bestandteile von $\mathcal{L}$ definiert. $\mathcal{L}$ wird auch als *die Prädikatenlogik erster Stufe für die Symbolmenge* $\{\in\}$ bezeichnet.

Für die Junktoren würden bereits z. B. $\neg$ und $\wedge$ genügen, denn $\vee$, $\rightarrow$, $\leftrightarrow$ lassen sich mit diesen Junktoren in ihrer gewünschten Bedeutung einführen. Ebenso ließe sich z. B. der Existenzquantor mit Hilfe des Allquantors einführen:

$$\begin{aligned}\varphi \vee \psi &= \neg(\neg\varphi \wedge \neg\psi), \\ \varphi \rightarrow \psi &= \neg\varphi \vee \psi, \\ \varphi \leftrightarrow \psi &= (\varphi \rightarrow \psi) \wedge (\psi \rightarrow \varphi), \\ \exists x\varphi &= \neg\forall x\neg\varphi.\end{aligned}$$

Beschränkt man sich in der Sprache auf $\neg$, $\wedge$ und $\forall$, so sind diese vier Gleichungen einfach Definitionen von $\vee$, $\rightarrow$, $\leftrightarrow$, $\exists$. In unserer reichhaltigeren Sprache sind $\varphi \vee \psi$ und $\neg(\neg\varphi \wedge \neg\psi)$ usw. verschiedene Formeln, die jedoch logisch äquivalent sind, d. h. wir können $\varphi \vee \psi$ genau dann formal beweisen, wenn wir $\neg(\neg\varphi \wedge \neg\psi)$ formal beweisen können. (Zum formalen Beweisbegriff s.u.)

Wir wollen Klammern wo immer möglich sparen, und vereinbaren hierzu die folgende Bindungsstärke der Quantoren und Junktoren: $\neg$, $\wedge$, $\vee$, $\rightarrow$, $\leftrightarrow$ (in absteigender Bindungsstärke).

Also meint $\varphi \wedge \psi \rightarrow \chi$ die Formel $(\varphi \wedge \psi) \rightarrow \chi$, und nicht $\varphi \wedge (\psi \rightarrow \chi)$.

Weiter schreiben wir $\forall x_1 \forall x_2 \dots \forall x_n \varphi$ kurz als $\forall x_1, \dots, x_n \varphi$. Ebenso für den Existenzquantor.

Unsere Sprache erlaubt keine Formeln der Gestalt $\forall x \in y \varphi$, $\exists x \in y \varphi$, wie sie in obiger Analyse von „ R ist eine Relation“ auftauchen. Wir führen hierzu zwei Abkürzungen ein, die sehr häufig verwendet werden:

$$\begin{aligned}\forall x \in y \varphi &= \forall x (x \in y \rightarrow \varphi), \\ \exists x \in y \varphi &= \exists x (x \in y \wedge \varphi).\end{aligned}$$

Die Punktnotation

Sehr viele Klammern kann man sich durch die *Punktnotation* sparen: Ein Punkt hinter einem Quantor bedeutet, daß der Wirkungsbereich eines Quantors so groß ist wie möglich. Z. B. ist

$\forall x. \varphi \rightarrow \exists y. \psi \wedge \chi$ die Formel $\forall x (\varphi \rightarrow \exists y (\psi \wedge \chi))$.

Weiter erlauben wir Klammern der Form $(\forall x. \dots)$ und $(\exists x. \dots)$.

So ist etwa $(\forall x. \varphi \rightarrow \exists y. \psi \wedge \chi) \rightarrow \rho$ die Formel $(\forall x (\varphi \rightarrow \exists y (\psi \wedge \chi)) \rightarrow \rho)$.

Diese Konventionen führen keine neuen Formeln ein, sondern erlauben lediglich, Formeln suggestiver und lesbarer zu notieren.

Freie Variablen und die Sätze von $\mathcal{L}$

Definition (*freie Variablen*)

Die freien Variablen einer Formel sind wie folgt festgelegt.

- (i) Die freien Variablen von $x = y$ und $x \in y$ sind x, y .
- (ii) Die freien Variablen von $\neg \varphi$ sind die freien Variablen von φ .
- (iii) Die freien Variablen von $\varphi \wedge \psi$, $\varphi \vee \psi$, $\varphi \rightarrow \psi$ und $\varphi \leftrightarrow \psi$ bestehen aus den freien Variablen von φ und den freien Variablen von ψ .
- (iv) Die freien Variablen von $\forall x \varphi$ und $\exists x \varphi$ sind die freien Variablen von φ ohne die Variable x .

Eine Variable x , die in φ in der Form $\forall x$ oder $\exists x$ vorkommt, heißt eine *gebundene Variable von φ* . Beispielsweise sind x, y, z die freien Variablen der Formel

$\varphi = (\exists u. u = y \vee u = x) \wedge \forall x. x \in y \rightarrow x = z$.

Die gebundenen Variablen dieser Formel sind u und x . x kommt also in φ sowohl frei als auch gebunden vor. Das erste Vorkommen von x in φ ist frei, das zweite Vorkommen von x in φ ist gebunden.

Definition (*Sätze/Aussagen*)

Eine Formel von $\mathcal{L}$ heißt eine *Aussage* oder ein *Satz* von $\mathcal{L}$, falls φ keine freien Variablen besitzt.

Definition (*Generalisierung und Allabschluß einer Formel*)

Sei φ eine Formel, und seien $x_1, \dots, x_n$ Variablen (mit $n \geq 0$).

Dann heißt $\forall x_1, \dots, x_n \varphi$ eine *Generalisierung von φ* .

Ist ψ eine Generalisierung von φ ohne freie Variablen, so heißt ψ ein *Allabschluß von φ* .

Der Fall $n = 0$ meint, daß φ selbst als eine Generalisierung von φ gilt.

Definition (die Schreibweise $\varphi(x_1, \dots, x_n)$)

Seien φ eine Formel und $x_1, \dots, x_n$ paarweise verschiedene Variablen.

Wir schreiben $\varphi(x_1, \dots, x_n)$, wenn die freien Variablen von φ unter $x_1, \dots, x_n$ vorkommen.

Schreiben wir $\varphi(x_1, \dots, x_n)$, so folgt nicht, daß genau $x_1, \dots, x_n$ die freien Variablen von φ sind. Z. B. können wir die Formel $\varphi = „\forall x \, x = y“$ schreiben als $\varphi(x, y, z)$. Diese Konvention erweist sich als überaus bequem.

Substitutionen

Sei $\varphi(x_1, x_2, \dots, x_n)$ eine Formel, und seien $y_1, \dots, y_n$ Variablen. Es gelte:

- (+) Setzen wir an der Stelle eines freien Vorkommens einer Variablen x_i in φ die Variable y_i statt x_i ein, so ist y_i wieder frei in φ .

y_i gerät dann durch einen solchen lokalen Austausch nicht in den Wirkungsbereich eines Quantors $\exists y_i$ oder $\forall y_i$ von φ .

Ist (+) erfüllt, so sei

$$\varphi_{x_1, \dots, x_n}(y_1, \dots, y_n)$$

die Formel, die entsteht, wenn wir in der Formel φ die freien Vorkommen der Variablen $x_1, \dots, x_n$ simultan durch $y_1, \dots, y_n$ ersetzen.

Ist (+) nicht erfüllt, so sei $\varphi_{x_1, \dots, x_n}(y_1, \dots, y_n)$ gleich φ .

Zuweilen schreiben wir auch kurz einfach $\varphi(y_1, \dots, y_n)$.

Sei etwa

$$\varphi(x) = \exists z \, z = x \wedge \forall y \, y = y \wedge \forall x \, x = x.$$

Dann ist (+) für $\varphi(x)$, y erfüllt und es gilt

$$\varphi_x(y) = \exists z \, z = y \wedge \forall y \, y = y \wedge \forall x \, x = x.$$

Dagegen ist (+) für $\varphi(x)$, z nicht erfüllt. Formal ergäbe ein Einsetzen von z in φ für x die Aussage $\exists z \, z = z \wedge \forall y \, y = y \wedge \forall x \, x = x$. Solche Variablenkollisionen durch Substitution vermeiden wir mit der Bedingung (+).

Mengentheoretische Eigenschaften

Wir können nun den im Aussonderungs- und Ersetzungsschema verwendeten Eigenschaftsbegriff festlegen, indem wir ihn einfach mit dem Formelbegriff identifizieren.

Definition (*mengentheoretische Eigenschaften und Aussagen*)

Eine *Eigenschaft* (der Mengenlehre) ist eine Formel von $\mathcal{L}$.

Eine *Aussage* (der Mengenlehre) ist ein Satz von $\mathcal{L}$.

Eine Eigenschaft ist also einfach ein aus den Grundrelationen „ $x = y$ “ und „ $x \in y$ “ mit Hilfe der Junktoren $\neg, \wedge, \vee, \rightarrow, \leftrightarrow$ und der Quantoren $\forall$ und $\exists$ aufgebauter Ausdruck. Und eine Aussage ist ein solcher Ausdruck, der keine freien Variablen enthält. Alle Sätze der beiden ersten Abschnitte ließen sich im Prinzip – bei geeigneter Definition von $\mathbb{N}, \mathbb{R}$, usw. – als Aussagen von $\mathcal{L}$ schreiben. Allerdings würden sie dadurch unlesbar.

Axiomensysteme

Definition (*Axiomensystem*)

Ein *Axiomensystem* von $\mathcal{L}$ ist eine endliche oder unendliche Liste

$\Sigma = \varphi_0, \varphi_1, \dots$ von Formeln von $\mathcal{L}$.

Implizit in dieser Festsetzung ist, daß ein Axiomensystem stets *rekursiv* ist, d.h. wir können von einer gegebenen Formel effektiv entscheiden, ob sie zum System gehört oder nicht.

Formale Beweise

Wir begnügen uns an dieser Stelle mit einer intuitiven und etwas vereinfachenden Beschreibung von formalen Beweisen. Der interessierte Leser findet am Ende dieses Kapitels eine detaillierte Beschreibung eines formalen Beweissystems, des sog. Hilbert-Kalküls.

Sei Σ ein Axiomensystem von $\mathcal{L}$. Ein *formaler Beweis* Δ aus Σ ist eine Liste

$\psi_0, \psi_1, \dots, \psi_m; \varphi_0, \varphi_1, \dots, \varphi_n$

von Formeln von $\mathcal{L}$, die nach bestimmten Regeln gebildet ist. $\psi_0, \dots, \psi_m$ sind dabei Axiome aus Σ . Intuition ist: Die Formelliste Δ ist ein Beweis von φ_n mit Hilfe der Axiome $\psi_0, \dots, \psi_m$. Die eigentliche Ableitung ist $\varphi_0, \dots, \varphi_n$. Hierbei entsteht ein φ_k durch Anwendung einer *Regel des Kalküls* auf Elemente von $\psi_0, \dots, \psi_m, \varphi_0, \dots, \varphi_{k-1}$.

Beispiele für Regeln sind:

(R1) Aus φ und ψ bilde $\varphi \wedge \psi$ (*Und-Einführung*),

(R2) Aus $\varphi \rightarrow \psi$ und φ bilde ψ (*Modus Ponens*).

Seien φ, ψ, χ drei beliebige Sätze von $\mathcal{L}$. Wir betrachten als Beispiel das Axiomensystem

$$\Sigma = \varphi, \psi, \varphi \wedge \psi \rightarrow \chi.$$

Wir können nun die Aussage χ aus Σ formal beweisen, denn

$$\Delta = \underset{\textcircled{1}}{\varphi}, \underset{\textcircled{2}}{\psi}, \underset{\textcircled{3}}{\varphi \wedge \psi} \rightarrow \underset{\textcircled{4}}{\chi}; \underset{\textcircled{5}}{\varphi \wedge \psi}, \underset{\textcircled{5}}{\chi}.$$

ist ein Beweis von χ mit Hilfe der Axiome aus Σ :

Die ersten drei Elemente $\textcircled{1}$, $\textcircled{2}$, $\textcircled{3}$ von Δ sind Axiome.

$\textcircled{4}$ entsteht durch Anwendung von (R1) auf $\textcircled{1}$ und $\textcircled{2}$.

$\textcircled{5}$, die durch Δ bewiesene Aussage, entsteht durch Anwendung von (R2) auf $\textcircled{3}$ und $\textcircled{4}$.

Im Prinzip ließe sich jeder umgangssprachlich bewiesene Satz der Mengenlehre in einer solchen Weise herleiten. Obwohl formale Beweise nur für sehr einfache Sätze der Mengenlehre oder allgemein der Mathematik direkt durchgeführt worden sind, ist man von der prinzipiellen Möglichkeit überzeugt, jeden in üblicher Weise bewiesenen Satz der Mathematik in einen formalen Beweis umwandeln zu können. Die Situation ist ähnlich zur prinzipiellen Auflösbarkeit einer Eigenschaft in ihre Bestandteile, ihre Umwandlung in eine Formel von $\mathcal{L}$: Schreiben wir einen üblichen, umgangssprachlichen mathematischen Beweis immer ausführlicher auf, so nähert sich seine Gestalt einem formalen Beweis. Umgangssprachliche Beweise von wenigen Zeilen können dann Hunderte von Regelanwendungen erforderlich machen – so wie die Auflösung von „R ist eine Relation“ schon eine relativ lange Formel ergibt.

Die ZFC-Axiome in der formalen Sprache

Wir formulieren nun das Axiomensystem ZFC in der Sprache der Mengenlehre. Vorab noch einige Bemerkungen zur Notation.

Hier und im folgenden sollen verschieden bezeichnete Variablen auch verschieden sein. So gilt dann z. B. für $\forall x \exists y \varphi(x, y)$ immer $x \neq y$ im Sinne verschiedener Variablen.

Der Quantor $\exists!$ bedeutet „es gibt genau ein“ und ist für jede Formel φ definiert durch

$$\exists! x \varphi = \exists y \forall x. \varphi \leftrightarrow x = y,$$

wobei y eine Variable ist, die in φ nicht vorkommt.

(EXT) Extensionalitätsaxiom

Zwei Mengen sind genau dann gleich, wenn sie die gleichen Elemente haben.

$$\forall x, y. x = y \leftrightarrow \forall z. z \in x \leftrightarrow z \in y$$

(LM) Existenz der leeren Menge

$$\exists x \forall y y \notin x$$

(PA) Paarmengenaxiom

Zu je zwei Mengen x, y existiert eine Menge z , die genau x und y als Elemente hat.

$$\forall x, y \exists z \forall u. u \in z \leftrightarrow u = x \vee u = y$$

(VER) Vereinigungsmengenaxiom

Zu jeder Menge x existiert eine Menge y , deren Elemente genau die Elemente der Elemente von x sind.

$$\forall x \exists y \forall z. z \in y \leftrightarrow \exists u. u \in x \wedge z \in u$$

(POT) Potenzmengenaxiom

Zu jeder Menge x existiert eine Menge y , die genau die Teilmengen von x als Elemente besitzt.

$$\forall x \exists y \forall z. z \in y \leftrightarrow \forall u. u \in z \rightarrow u \in x$$

(AUS) Aussonderungsschema

Zu jeder Eigenschaft φ und jeder Menge x gibt es eine Menge y , die genau die Elemente von x enthält, auf die φ zutrifft.

$$\forall x, p_1, \dots, p_n \exists y \forall u. u \in y \leftrightarrow u \in x \wedge \varphi(u, p_1, \dots, p_n)$$

(ERS) Ersetzungsschema

Das Bild einer Menge unter einer Funktion φ ist eine Menge.

$$\forall p_1, \dots, p_n. \forall u \exists! v \varphi(u, v, p_1, \dots, p_n) \rightarrow \\ \forall x \exists y \forall v. v \in y \leftrightarrow \exists u. u \in x \wedge \varphi(u, v, p_1, \dots, p_n)$$

(UN) Unendlichkeitsaxiom

Es existiert eine Menge x , die die leere Menge als Element enthält und die mit jedem ihrer Elemente y auch $\{y\}$ als Element enthält.

$$\exists x. (\exists y. y \in x \wedge \forall z z \notin y) \wedge \forall y. y \in x \rightarrow \exists z. z \in x \wedge \forall u. u \in z \leftrightarrow u = y$$

(FUN) Fundierungsaxiom

Jede nichtleere Menge x hat ein Element y , das mit x kein Element gemeinsam hat.

$$\forall x. \exists y y \in x \rightarrow \exists y. y \in x \wedge \forall z. z \in y \rightarrow z \notin x$$

(AC) Auswahlaxiom

*Ist x eine Menge, deren Elemente nichtleer und paarweise disjunkt sind,
so existiert eine Menge y ,
die mit jedem Element von x genau ein Element gemeinsam hat.*

$$\forall x. (\forall y (y \in x \rightarrow \exists z z \in y) \wedge \forall y, z. y \in x \wedge z \in x \wedge y \neq z \rightarrow \forall u. u \in y \rightarrow u \notin z) \\ \rightarrow \exists y \forall z. z \in x \rightarrow \exists! u. u \in y \wedge u \in z$$

Das Axiomensystem von $\mathcal{L}$ bestehend aus den unendlich vielen angegebenen Aussagen von $\mathcal{L}$ bezeichnen wir wie das umgangssprachliche System ebenfalls mit ZFC.

Formale Beweise im Hilbert-Kalkül

Wir besprechen nun ausführlicher einen formalen Beweisbegriff für die Sprache $\mathcal{L}$ der Mengenlehre. Damit wird die Formalisierung – oder genauer: die prinzipielle Möglichkeit der Formalisierung – der Mengenlehre vollständig.

Ansätze zu formalen Beweissystemen finden sich in der Syllogistik des Aristoteles. Die Logik blieb dann auch viele Jahrhunderte unter dem Dach der Philosophie hängen, bis sie von der Mathematik als Objekt der Begierde erkannt wurde und in Folge dann endlich den Auszug aus dem Elternhaus hinter sich brachte. Den wichtigsten Fortschritt nach Aristoteles bildeten die Arbeiten von Gottlob Frege (1848 – 1925), und in den Jahrzehnten nach 1900 wurden dann die ersten brauchbaren mathematischen Logikkalküle entwickelt. Prominent sind heute der (Frege-) Hilbert-Kalkül und der Kalkül des natürlichen Schließens von Gerhard Gentzen. Wir verwenden hier den Hilbert-Kalkül, der sich einfach definieren läßt, und der darüber hinaus zur Gewinnung von Resultaten der mathematischen Logik überaus gut zu handhaben ist.

Für das formale Beweisen selbst ist der Kalkül des natürlichen Schließens angenehmer, insbesondere in seiner zweidimensionalen Form, bei der Beweise in Form von Bäumen notiert werden. Der Leser konsultiere Lehrbücher zur mathematischen Logik für die Feinheiten der verschiedenen Kalküle. Die verschiedenen Kalküle sind aber allesamt logisch äquivalent, d.h. sie beweisen dieselben Aussagen. Die Beweise der Kalküle lassen sich darüber hinaus effektiv ineinander übersetzen.

Der Hilbert-Kalkül für $\mathcal{L}$ ist ein syntaktisches System auf der Metaebene. Es beschreibt, welche endlichen Folgen von Formeln einen formalen Beweis bilden. Ein formaler Beweis beweist dann immer seine letzte Formel.

Der Hilbert-Kalkül hat zwei Bestandteile: die logischen Axiome und die Schlußregel Modus Ponens. Die logischen Axiome zerfallen in aussagenlogische Axiome, Identitätsaxiome und Quantoraxiome.

Die logischen Axiome des Hilbert-Kalküls

Wir zeichnen bestimmte Formeln von $\mathcal{L}$ als logische Grundaxiome aus. Man kann diese Formeln als stillschweigende Axiome ansehen, die jedes Axiomensystem schon mitbringt. Sie liefern keinen mathematischen Inhalt, sondern verbinden die logischen Zeichen miteinander, getreu der intendierten Semantik.

Im Hilbert-Kalkül tauchen wegen des ständigen Einsatzes der Modus-Ponens-Schlußregel (s. u.) sehr häufig Implikationsketten auf. Wir vereinbaren hier Rechtsklammerung, um Klammern zu sparen: $\varphi \rightarrow \psi \rightarrow \chi$ ist also $\varphi \rightarrow (\psi \rightarrow \chi)$, und $\varphi \rightarrow \psi \rightarrow \chi \rightarrow \xi$ ist $\varphi \rightarrow (\psi \rightarrow (\chi \rightarrow \xi))$, usw. Hilfreich für das Lesen solcher Ketten ist die semantische Äquivalenz von „aus φ folgt: aus ψ folgt χ “ mit „aus φ und ψ folgt χ “. Allgemein ist „ φ_1 folgt φ_2 folgt ... folgt φ_n folgt φ_{n+1} “ mit Rechtsklammerung äquivalent zu „aus φ_1 und ... und φ_n folgt φ_{n+1} “, wobei hier die φ_i beliebige umgangssprachliche mathematische Aussagen sein sollen. Mit dieser Lesart wird etwa (i) und (ii) in der folgenden Definition sofort einsichtig.

Definition (aussagenlogische Axiome)

Eine Formel von $\mathcal{L}$ heißt ein *aussagenlogisches Axiom* von $\mathcal{L}$, wenn sie von einer der folgenden Formen ist:

- (i) $\varphi \rightarrow \psi \rightarrow \varphi$,
- (ii) $(\varphi \rightarrow \psi \rightarrow \chi) \rightarrow (\varphi \rightarrow \psi) \rightarrow \varphi \rightarrow \chi$,
- (iii) $\varphi \wedge \psi \rightarrow \varphi$,
- (iv) $\varphi \wedge \psi \rightarrow \psi$,
- (v) $\varphi \rightarrow \psi \rightarrow \varphi \wedge \psi$,
- (vi) $\varphi \rightarrow \varphi \vee \psi$,
- (vii) $\varphi \rightarrow \psi \vee \varphi$,
- (viii) $(\varphi \rightarrow \chi) \rightarrow (\psi \rightarrow \chi) \rightarrow \varphi \vee \psi \rightarrow \chi$,
- (ix) $(\varphi \rightarrow \psi) \rightarrow (\psi \rightarrow \varphi) \rightarrow (\varphi \leftrightarrow \psi)$,
- (x) $(\varphi \leftrightarrow \psi) \rightarrow (\varphi \rightarrow \psi) \wedge (\psi \rightarrow \varphi)$,
- (xi) $\varphi \wedge \neg \varphi \rightarrow \psi$,
- (xii) $\neg \neg \varphi \rightarrow \varphi$.

Man erhält insgesamt einen logisch äquivalenten Kalkül, wenn man hier liberaler jede aussagenlogische Tautologie als Axiom zuläßt. Eine aussagenlogische Tautologie ist dabei eine Formel von $\mathcal{L}$, die, quantorfrei notiert, durch das bekannte Wahrheitstafelverfahren für lange Winterabende als wahr in jedem Falle nachgewiesen werden kann. Wir können zum Beispiel $\forall x \varphi \rightarrow \forall x \varphi \vee \exists y \varphi$ in der quantorfreien Form $\psi \rightarrow \psi \vee \chi$ notieren, und alle vier Kombinationen von Wahrheitswerten w (für wahr) oder f (für falsch) für ψ und χ durchspielen, um zu sehen, daß die Auswertung von $\vee$ und $\rightarrow$ mit der üblichen Semantik am Ende in allen vier Fällen w ergibt. Es gilt etwa $f \vee f = f$, $f \rightarrow f = w$, also $f \rightarrow (f \vee f) = f \rightarrow f = w$, ebenso für die anderen drei Fälle. Allgemeiner hat man 2^n -viele Kombinationen zu überprüfen, wenn eine quantorfreie Formel aus $\varphi_1, \dots, \varphi_n$ zusammengesetzt ist.

Obige Liste von Schemata ist aber übersichtlicher und interessanter als ein allgemeines Tautologie-Schema (alle Axiome der Liste sind Tautologien). Mit ihrer Hilfe kann man zudem sehr leicht sinnvolle restriktive Logiken genau definieren: Der einzige Stein, an dem sich die konstruktive Mathematik im ganzen Hilbert-Park das Knie schlägt, ist das Axiom (xii). Dieses *Stabilitäts-* oder *reductio ad absurdum*-Schema $\neg\neg\varphi \rightarrow \varphi$ macht die Logik zur klassischen Logik. Läßt man (xii) weg, so erhält man den Hilbert-Kalkül für die – für den klassisch geschulten Geist nicht unbedingt intuitive – *intuitionistische Logik* (befürwortet vor allem von Brouwer, und formalisiert von Heyting (1930); vgl. auch [Kolmogorov 1925]). Intuitionistisch ist dagegen spaßigerweise immer noch $\neg\neg\neg\varphi \rightarrow \neg\varphi$ beweisbar, von drei Negationen kann man also zwei fallen lassen.

Gleichwertig zu (xii) über den anderen logischen Axiomen ist das *tertium non datur*-Schema $\varphi \vee \neg\varphi$, das die konstruktive Mathematik ebenso scheut wie das Wegstreichen einer doppelten Negation.

Die sog. Brouwer-Heyting-Kolmogorov Interpretation klärt, was die intuitionistische Logik will: Unter dieser (informalen) Interpretation gilt: Ein Beweis von $\varphi \wedge \psi$ liegt genau dann vor, wenn ein Beweis von φ und ein Beweis von ψ vorliegt. (Das gilt klassisch genauso.) Ein Beweis von $\varphi \vee \psi$ liegt genau dann vor, wenn ein Beweis von φ oder ein Beweis von ψ vorliegt. (Das gilt klassisch nicht: klassisch ist z. B. $\varphi \vee \neg\varphi$ für jede Aussage φ beweisbar, während für jede offene Frage oder jede unabhängige Aussage φ weder ein Beweis von φ noch ein Beweis von $\neg\varphi$ vorliegt.) Weiter liegt ein Beweis von $\varphi \rightarrow \psi$ genau dann vor, wenn es ein Verfahren gibt, das einen beliebigen Beweis von φ in einen Beweis von ψ transformiert. Und schließlich liegt ein Beweis von $\neg\varphi$ genau dann vor, wenn es ein Verfahren gibt, das einen Beweis von φ in einen Beweis einer Formel verwandeln würde, deren Nichtbeweisbarkeit man akzeptiert. (Dies gilt klassisch ebenfalls nicht. Das Schema (xii) $\neg\neg\varphi \rightarrow \varphi$ ist gerade etwas, was man mit einer reinen Beweistransformation nicht begründen kann, im Gegensatz zu $\neg\neg\neg\varphi \rightarrow \neg\varphi$ oder auch $\varphi \rightarrow \neg\neg\varphi$.)

$\neg\neg\varphi \rightarrow \varphi$ wird in der Mathematik durchgehend verwendet, und die Art und Weise der Verwendung erklärt den Namen *reductio ad absurdum*: Man will φ zeigen. Hierzu macht man die Annahme $\text{non}(\varphi)$, und zeigt, daß aus dieser Annahme Unfug wie „ $0 = 1$ “ hergeleitet werden kann. Dann hat man gezeigt, daß $\text{non}(\varphi)$ nicht gelten kann, d. h. man hat durch das Argument $\text{non}(\text{non}(\varphi))$ gezeigt. Bis hierhin spielt der Konstruktivist mit (bei konstruktiver Argumentation), nicht aber bei dem letzten Schluß von $\text{non}(\text{non}(\varphi))$ auf φ . Man muß den Konstruktivisten rechtgeben, daß ein ständiges „Annahme $\text{non}(\varphi)$ “ zur Manie werden kann, und häufig auch dort eingesetzt wird, wo es nicht nötig ist. Ob das schlimm ist, ist die andere Frage.

Verwandt zu *reductio ad absurdum* ist das *Kontrapositionsgesetz* oder das Schema des *indirekten Beweisens* $\varphi \rightarrow \psi \leftrightarrow \neg\psi \rightarrow \neg\varphi$. Intuitionistisch gilt $(\varphi \rightarrow \psi) \rightarrow \neg\psi \rightarrow \neg\varphi$, i. a. aber nicht die Umkehrung, die für das indirekte Beweisen gebraucht wird.

Ein noch restriktiveres System ist die *Minimallogik* [Johansson 1937]. Hier läßt man neben (xii) auch noch das *ex contradictio quodlibet*-Schema (xi) $\varphi \wedge \neg\varphi \rightarrow \psi$ weg. Zudem streicht man das Negationszeichen $\neg$ aus der Sprache, führt einen 0-stelligen logischen Junktor $\perp$, *falsum* genannt, ein, und definiert dann $\neg\varphi$ als die Formel $\varphi \rightarrow \perp$. Das Zeichen $\perp$ selbst gilt als Primformel. Dieses Vorgehen ist auch im Intuitionismus üblich, und $\neg\varphi$ als $\varphi \rightarrow \perp$ zu lesen deckt sich mit der obigen Interpretation, sobald $\perp$ als eine Formel ohne Beweis angesehen wird. Klassisch kann man den unbestreitbaren formalen Komfort eines Falsums erreichen, indem man $\perp$ als Abkürzung für $\psi \wedge \neg\psi$ einführt, mit einem beliebigen ψ . Damit ist die nützliche Äquivalenz $\neg\varphi \leftrightarrow \varphi \rightarrow \perp$ für alle φ formal beweisbar.

In grober Analogie: Die klassische Logik verhält sich zur intuitionistischen Logik so wie ZFC zu ZF. Und so wie ZF eine interessante Teiltheorie von ZFC ist, ist eine Logik ohne Stabilität zwar etwas wackelig, aber nicht uninteressant.

Definition (*Identitätsaxiome*)

Eine Formel von $\mathcal{L}$ heißt ein *Identitätsaxiom* von $\mathcal{L}$, wenn sie von einer der folgenden Formen ist:

- (i) $x = x$,
- (ii) $x = y \rightarrow \varphi_z(x) \rightarrow \varphi_z(y)$, wobei φ eine Primformel ist.

Definition (*Quantoraxiome*)

Eine Formel von $\mathcal{L}$ heißt ein *Quantoraxiom* von $\mathcal{L}$, wenn sie von einer der folgenden Formen ist:

- (i) $\varphi \rightarrow \forall x \varphi$, wobei x nicht frei in φ vorkommt,
- (ii) $\forall x(\varphi \rightarrow \psi) \rightarrow \forall x \varphi \rightarrow \forall x \psi$,
- (iii) $\forall x \varphi \rightarrow \varphi_x(y)$,
- (iv) $\varphi_x(y) \rightarrow \exists x \varphi$,
- (v) $\forall x(\varphi \rightarrow \psi) \rightarrow \exists x \varphi \rightarrow \psi$, wobei x nicht frei in ψ vorkommt.

Die an dieser Stelle etwas seltsame Variablenbedingung im ersten Schema wird unten klar werden. Semantisch sind die fünf Schemata sicher korrekt.

Definition (*logische Axiome des Hilbert-Kalküls für $\mathcal{L}$*)

Eine Formel φ von $\mathcal{L}$ heißt ein *logisches Axiom* von $\mathcal{L}$, falls φ eine Generalisierung eines aussagenlogischen Axioms, eines Identitätsaxioms oder eines Quantoraxioms ist.

Ist also φ ein logisches Axiom, so ist $\forall x \varphi$ ein logisches Axiom für alle Variablen x . Eine Formel φ gilt als Generalisierung von φ , und damit sind alle aufgelisteten Axiome auch logische Axiome. Die Motivation für die technische Variablenbehandlung in der Definition der logischen Axiome wird aus dem Beweis der Generalisierungsregel hervorgehen (s. u.).

Die Schlußregel des Hilbert-Kalküls

Im Hilbert-Kalkül gibt es nur eine Schlußregel, den Modus Ponens:

Definition (*Modus Ponens*)

Seien φ, ψ, χ Formeln von $\mathcal{L}$. χ entsteht aus φ, ψ durch Anwendung der Regel *Modus Ponens*, falls ψ von der Form $\varphi \rightarrow \chi$ ist.

Die Idee ist: Haben wir φ und $\varphi \rightarrow \chi$ bewiesen, so können wir auf χ schließen. Und diesen Schluß nennen wir Modus Ponens.

Formale Beweise im Hilbert-Kalkül

Definition (*formale Beweise im Hilbert-Kalkül für die Sprache der Mengenlehre*)

Sei Σ ein Axiomensystem von $\mathcal{L}$.

Weiter sei $\Delta = \varphi_1, \dots, \varphi_n$ eine endliche Liste von $\mathcal{L}$ -Formeln.

Δ heißt *ein (formaler) Beweis aus Σ* oder *eine Herleitung aus Σ* ,

falls für alle $1 \leq i \leq n$ gilt:

- (i) φ_i ist ein logisches Axiom oder
- (ii) $\varphi_i \in \Sigma$ oder
- (iii) es existieren $1 \leq k, \ell < i$ derart, daß φ_i aus φ_k, φ_ℓ durch Anwendung von Modus Ponens entsteht.

Definition (*Beweisbarkeit einer Formel, Länge eines Beweises*)

Sei Σ ein Axiomensystem von $\mathcal{L}$, und sei φ eine Formel von $\mathcal{L}$.

φ heißt (*formal*) *beweisbar aus Σ* oder *herleitbar aus Σ* , in Zeichen

$\Sigma \vdash \varphi$ [Σ beweist φ], falls gilt:

Es existiert ein Beweis $\Delta = \varphi_1, \dots, \varphi_n$ aus Σ mit $\varphi = \varphi_n$.

Δ heißt dann *ein Beweis von φ aus Σ der Länge n* .

Wir schreiben $\vdash \varphi$, falls ein leeres Axiomensystem φ beweist, d.h.

falls ein Beweis von φ existiert, der nur (i) und (iii) verwendet.

Gilt $\vdash \varphi$, so nennen wir φ (*logisch*) *beweisbar*.

Analog sind $\Sigma, \Sigma' \vdash \varphi$ sowie $\Sigma, \psi \vdash \varphi$ definiert für Axiomensysteme Σ, Σ' und Formeln $\psi, \Sigma, \Sigma' \vdash \varphi$ heißt, daß in einem Beweis von φ Axiome aus Σ und Σ' verwendet werden dürfen. $\Sigma, \psi \vdash \varphi$ heißt, daß in einem Beweis von φ neben Axiomen aus Σ auch die Formel ψ als Axiom verwendet werden darf.

Es gilt $\Sigma \vdash \varphi$ (in einem Schritt) für alle $\varphi \in \Sigma$ und $\varphi \vdash \varphi$ für alle Formeln φ .

Durchgehend verwendet werden die folgenden Beobachtungen einfacher Natur:

Übung

Für alle Axiomensysteme Σ und alle Formeln φ gilt:

- (i) Ist $\Delta = \varphi_1, \dots, \varphi_n$ ein Beweis aus Σ , so gilt $\Sigma \vdash \varphi_i$ für alle $1 \leq i \leq n$.
- (ii) Gilt $\Sigma \vdash \varphi$, so existieren Axiome $\psi_1, \dots, \psi_n$ aus Σ mit $\psi_1, \dots, \psi_n \vdash \varphi$.

Anfangsstücke von Beweisen sind Beweise. Und formale Beweise sind immer finit: Obwohl ein Axiomensystem aus unendlich vielen Formeln bestehen kann, werden in einem Beweis immer nur endlich viele Axiome verwendet.

In jedem Kalkül verwenden Beweise Grundannahmen, und ziehen Schlüsse aus diesen Annahmen. Die eigentliche Dynamik eines Beweises im Hilbert-Kalkül steckt im Modus Ponens. Wenn wir in nichttrivialer Weise χ beweisen wollen, müssen wir vorher irgendwann für ein φ die Formeln $\varphi \rightarrow \chi$ und φ bewiesen haben. Welches φ man hierbei anstrebt, ist die Kunst. Das gleiche gilt nun aber für φ : φ ist entweder ein Axiom, oder entstanden aus $\psi \rightarrow \varphi$ und ψ . Ähnliches gilt für ψ . Trivial oder Modus Ponens. Tertium non datur.

Das tatsächliche formale Beweisen im Hilbert-Kalkül ist qualvoll. Dante hätte es in seiner Göttlichen Komödie gut verwenden können. Wer während seines Lebens etwas gegen die Logik gesagt hat, muß in der Hölle formale Beweise im Hilbert-Kalkül führen. Auf Erden verschafft der folgende Satz aus der logischen Apotheke Linderung:

Satz (*Deduktionstheorem*)

Seien Σ ein Axiomensystem, und seien ψ, φ Formeln von $\mathcal{L}$.

Dann sind äquivalent:

- (i) $\Sigma, \psi \vdash \varphi$,
- (ii) $\Sigma \vdash \psi \rightarrow \varphi$.

Beweis

(i) $\hookrightarrow$ (ii):

Wir zeigen die Aussage durch metamathematische Induktion über die Länge $n \geq 1$ eines formalen Beweises Δ von φ aus Σ, ψ . Entscheidend hierbei sind die aussagenlogischen Axiome (i) und (ii).

Induktionsschritt n

Sei $\Delta = \varphi_1, \varphi_2, \dots, \varphi_n$ ein Beweis von $\varphi_n = \varphi$ aus Σ, ψ .

1. *Fall*: φ ist ein logisches Axiom oder ein Axiom von Σ .

Dann ist $\varphi, \varphi \rightarrow \psi \rightarrow \varphi, \psi \rightarrow \varphi$ ein Beweis von $\psi \rightarrow \varphi$ aus Σ .

2. *Fall*: φ entsteht durch Anwendung von Modus Ponens.

Dann existieren $k, \ell < n$ mit $\varphi_\ell = \varphi_k \rightarrow \varphi$.

Nach Induktionsvoraussetzung existieren Beweise

- (a) $\eta_1, \dots, \eta_m$ von $\eta_m = \psi \rightarrow \varphi_\ell = \psi \rightarrow \varphi_k \rightarrow \varphi$ aus Σ ,
- (b) $\xi_1, \dots, \xi_{m'}$ von $\xi_{m'} = \psi \rightarrow \varphi_k$ aus Σ .

Dann ist

$\eta_1, \dots, \eta_{m-1}, \psi \rightarrow \varphi_k \rightarrow \varphi, \xi_1, \dots, \xi_{m'-1}, \psi \rightarrow \varphi_k,$
 $(\psi \rightarrow \varphi_k \rightarrow \varphi) \rightarrow (\psi \rightarrow \varphi_k) \rightarrow \psi \rightarrow \varphi, (\psi \rightarrow \varphi_k) \rightarrow \psi \rightarrow \varphi, \psi \rightarrow \varphi$
 ein Beweis von $\psi \rightarrow \varphi$ aus Σ .

(ii) $\hookrightarrow$ (i):

Sei $\varphi_1, \dots, \varphi_n$ ein Beweis von $\varphi_n = \psi \rightarrow \varphi$ aus Σ .

— Dann ist $\varphi_1, \dots, \varphi_{n-1}, \psi \rightarrow \varphi, \psi, \varphi$ ein Beweis von φ aus Σ, ψ .

Die absolutistische Existenz des Modus Ponens ist eine Besonderheit des Hilbert-Kalküls. Diese Regierungsform wird zwar von einer Unzahl logischer Axiome gestützt, aber ohne das Deduktionstheorem käme es sofort zur Revolution. Der Leser wird bei den Übungen unten sehen, wie nützlich das Deduktionstheorem tatsächlich ist. Hier ein einfaches Beispiel: Wir zeigen

$\varphi \rightarrow \psi \rightarrow \chi \vdash \psi \rightarrow \varphi \rightarrow \chi$.

Denn $\varphi \rightarrow \psi \rightarrow \chi \vdash \varphi \rightarrow \psi \rightarrow \chi$. Zweimalige Anwendung des Deduktionstheorems liefert $\varphi \rightarrow \psi \rightarrow \chi, \varphi, \psi \vdash \chi$. Trivialerweise also $\varphi \rightarrow \psi \rightarrow \chi, \psi, \varphi \vdash \chi$. Zweimalige Anwendung der einfachen Richtung des Deduktionstheorems liefert nun $\varphi \rightarrow \psi \rightarrow \chi \vdash \psi \rightarrow \varphi \rightarrow \chi$.

Daneben ist der Satz von theoretischer Bedeutung. Hierzu betrachten wir noch eine Verallgemeinerung. Durch Induktion folgt aus dem Deduktionstheorem, daß für alle Axiomensysteme Σ und alle Formeln $\varphi, \psi_1, \dots, \psi_n$ gilt:

$$\Sigma, \psi_1, \dots, \psi_n \vdash \varphi \text{ gdw } \Sigma \vdash \psi_1 \rightarrow \dots \rightarrow \psi_n \rightarrow \varphi.$$

Es folgt, daß alles, was ein endliches Axiomensystem beweisen kann, sich im wesentlichen in der reinen Logik abspielt. Denn ist $\Sigma = \psi_1, \dots, \psi_n$, so gilt für alle φ :

$$\Sigma \vdash \varphi \text{ gdw } \vdash \psi_1 \rightarrow \dots \rightarrow \psi_n \rightarrow \varphi.$$

Ein zum Beweis des Deduktionstheorems analoges Argument liefert die folgende Generalisierungsregel:

Satz (*Generalisierungsregel*)

Sei Σ ein Axiomensystem von $\mathcal{L}$, und sei φ eine Formel von $\mathcal{L}$ mit $\Sigma \vdash \varphi$.

Weiter sei x eine Variable, die in keiner Formel von Σ frei vorkommt.

Dann gilt $\Sigma \vdash \forall x \varphi$.

Beweis

Wir zeigen die Aussage durch metamathematische Induktion über die Länge $n \geq 1$ eines formalen Beweises Δ von φ aus Σ .

Entscheidend hierbei sind die Quantoraxiome (i) und (ii).

Induktionsschritt n

Sei $\Delta = \varphi_1, \varphi_2, \dots, \varphi_n$ ein Beweis von $\varphi_n = \varphi$ aus Σ .

1. *Fall*: φ ist ein logisches Axiom oder ein Element von Σ .

Dann ist $\varphi, \varphi \rightarrow \forall x \varphi, \forall x \varphi$ ein Beweis von $\forall x \varphi$ aus Σ , denn nach Voraussetzung ist x nicht frei in φ_n .

2. *Fall*: φ entsteht durch Anwendung von Modus Ponens.

Dann existieren $k, \ell < n$ mit $\varphi_\ell = \varphi_k \rightarrow \varphi$.

Nach Induktionsvoraussetzung existieren Beweise

(a) $\eta_1, \dots, \eta_m$ von $\eta_m = \forall x \varphi_\ell = \forall x. \varphi_k \rightarrow \varphi$ aus Σ ,

(b) $\xi_1, \dots, \xi_{m'}$ von $\xi_{m'} = \forall x \varphi_k$ aus Σ .

Dann ist

$$\eta_1, \dots, \eta_{m-1}, \forall x. \varphi_k \rightarrow \varphi, \xi_1, \dots, \xi_{m'-1}, \forall x \varphi_k, \\ (\forall x. \varphi_k \rightarrow \varphi) \rightarrow \forall x \varphi_k \rightarrow \forall x \varphi, \forall x \varphi_k \rightarrow \forall x \varphi, \forall x \varphi$$

— ein Beweis von $\forall x \varphi$ aus Σ .

Besteht ein Axiomensystem Σ , wie etwa ZFC, nur aus Aussagen, so beweist Σ eine Formel $\varphi(x_1, \dots, x_n)$ genau dann, wenn Σ die Aussage $\forall x_1, \dots, x_n \varphi$ beweist. (Die Richtung von rechts nach links erhält man durch iterierte Anwendung des Quantoraxioms $\forall x \psi \rightarrow \psi_x(y)$ und Modus Ponens.)

In manchen Formulierungen des Hilbert-Kalküls wird eine Generalisierungsregel als zweite Schlußregel neben dem Modus Ponens verwendet: Von φ darf man immer auf die Formel $\forall x \varphi$ schließen, wenn x nicht frei in φ vorkommt. Ein Axiomensystem Σ muß dann ausschließlich aus Aussagen bestehen, was keine wesentliche Einschränkung ist. Eine einzige Schlußregel zu haben wird aber dem Hilbert-Kalkül als Staatsform sicher gerechter und vereinfacht seine metatheoretische Untersuchung.

Übung

Für alle Formeln φ, ψ, χ , alle Variablen x und alle Axiomensysteme Σ gilt:

- (i) $\vdash \varphi \rightarrow \psi \rightarrow \chi \leftrightarrow \varphi \wedge \psi \rightarrow \chi$,
- (ii) $\vdash \neg \varphi \leftrightarrow \varphi \rightarrow \psi \wedge \neg \psi$,
- (iii) $\vdash \varphi \vee \neg \varphi$,
- (iv) $\Sigma, \varphi \vdash \psi$ und $\Sigma, \neg \varphi \vdash \psi$ folgt $\Sigma \vdash \psi$,
- (v) $\vdash (\varphi \rightarrow \neg \varphi) \rightarrow \neg \varphi$,
- (vi) $\vdash (\varphi \vee \psi) \rightarrow \neg \psi \rightarrow \varphi$,
- (vii) $\vdash \varphi \rightarrow \psi \leftrightarrow \neg \varphi \vee \psi$,
- (viii) $\vdash \varphi \rightarrow \psi \leftrightarrow \neg \psi \rightarrow \neg \varphi$,
- (ix) $\vdash \varphi \vee \psi \leftrightarrow \neg (\neg \varphi \wedge \neg \psi)$,
- (x) $\vdash \forall x \varphi \leftrightarrow \neg \exists x \neg \varphi$.
- (xi) $\vdash \exists x \varphi \leftrightarrow \neg \forall x \neg \varphi$.

Es ist instruktiv zu verfolgen, für welche formale Beweise hierbei (und für welche Richtungen von $\leftrightarrow$) das Schema $\neg \neg \varphi \rightarrow \varphi$ verwendet wird.

Widerspruchsfreiheit

Definition (Widerspruchsfreiheit)

Sei Σ ein Axiomensystem von $\mathcal{L}$. Σ heißt *widerspruchsfrei* oder *konsistent*, falls es eine Formel φ gibt mit $\text{non}(\Sigma \vdash \varphi)$.

Nach der Übung oben gilt: Σ ist widerspruchsfrei *gdw* jeder endliche Teil Σ' von Σ ist widerspruchsfrei.

Wir halten eine fundamentale Äquivalenz fest.

Satz (widerspruchsfreie Erweiterungen)

Seien Σ ein Axiomensystem und φ eine Formel von $\mathcal{L}$.

Dann sind äquivalent:

- (i) Σ erweitert um das Axiom φ ist widerspruchsfrei.
- (ii) $\text{non}(\Sigma \vdash \neg \varphi)$.

Beweis

(i) $\hookrightarrow$ (ii): Sei ψ eine Formel mit $\text{non}(\Sigma, \varphi \vdash \psi)$.

Annahme, es gilt $\Sigma \vdash \neg \varphi$.

Sei dann $\varphi_1, \dots, \varphi_n = \neg \varphi$ ein Beweis von $\neg \varphi$ aus Σ .

Dann ist aber

$\varphi_1, \dots, \varphi_{n-1}, \neg \varphi, \varphi, \neg \varphi \rightarrow \varphi \rightarrow \psi, \varphi \rightarrow \psi, \psi$,

ein Beweis von ψ aus Σ, φ , wobei wir die beweisbare Formel

$\neg \varphi \rightarrow \varphi \rightarrow \psi$ verwenden.

Dies ist ein *Widerspruch* zur Wahl von ψ .

(ii) $\hookrightarrow$ (i): Wir zeigen $\text{non}(i) \hookrightarrow \text{non}(ii)$.

Sei also Σ erweitert um φ widerspruchsvoll.

Dann gilt $\Sigma, \varphi \vdash \neg \varphi$.

Nach dem Deduktionstheorem gilt also $\Sigma \vdash \varphi \rightarrow \neg \varphi$.

Aber $(\varphi \rightarrow \neg \varphi) \rightarrow \neg \varphi$ ist beweisbar.

— Modus Ponens liefert dann $\Sigma \vdash \neg \varphi$, also $\text{non}(ii)$.

Aus $\text{non}(\Sigma \vdash \varphi)$ kann i.a. nicht auf $\Sigma \vdash \neg \varphi$ geschlossen werden. Die Existenz unabhängiger Aussagen macht diesen Schluß zunichte.

Definition (*unabhängig*)

Sei Σ ein Axiomensystem, und sei φ eine Formel von $\mathcal{L}$.

φ heißt *unabhängig von Σ* , falls gilt:

$\text{non}(\Sigma \vdash \varphi)$ und $\text{non}(\Sigma \vdash \neg \varphi)$.

Die Existenz einer von Σ unabhängigen Aussage impliziert offenbar die Widerspruchsfreiheit von Σ . Andererseits besagt der erste Gödelsche Unvollständigkeitssatz: Ist eine axiomatische Erweiterung Σ von ZFC widerspruchsfrei, so existieren immer von Σ unabhängige Aussagen. Also decken auch beliebig starke Erweiterungen von ZFC nie alles ab. Die unabhängigen Aussagen des Gödelresultats sagen etwa „ich bin nicht beweisbar in Σ “. Daß solche Aussagen unabhängig sind, irritiert nicht unbedingt. Das wirklich Aufregende ist nun aber, daß wir über ZFC unabhängige Aussagen mit einem klaren mengentheoretischen Gehalt kennen:

Formale Form des Fundamentalsatzes der Mengenlehre

Ist ZFC widerspruchsfrei, so ist (CH) unabhängig von ZFC.

Dieser Satz ist (wie etwa das Deduktionstheorem) ein Metatheorem. Für den Beweis wird viel Mengenlehre verwendet, und keine Syntaxanalyse des Kalküls. Letztendlich liefert der Beweis des Fundamentalsatzes dann aber sogar ein Verfahren, das zeigt, wie wir einen Beweis von (CH) oder von $\neg$ (CH) in ZFC in einen Beweis eines Widerspruchs $\psi \wedge \neg \psi$ verwandeln können.

Übersetzung des formalen Rahmens in die Objekzebene

Man kann mathematische Logik ganz allgemein in der Mengenlehre entwickeln, und $\mathcal{L}$ erscheint dann in der Objekzebene als spezielle prädikatenlogische Sprache. Als Zeichen der Sprache wählt man geeignete Mengen, etwa $\wedge = 3$, $\forall = 7$, $v_0 = 0$, $v_1 = 2$, $v_2 = 4$, usw. Formeln sind dann wieder bestimmte Folgen von Zeichen, und formale Beweise auf der Objekzebene sind bestimmte endliche Folgen von Formeln. Der Transport von Begriffen und Resultaten funktioniert problemlos von der Metaebene in die Objekzebene, und wird üblicherweise durch Klammern $\lceil \cdot \rceil$ bezeichnet. Etwa $\lceil \wedge \rceil = 3$, $\lceil \forall v_2 \wedge \wedge \rceil = (7, 4, 3, 3)$. Formale Beweise der wirklichen Welt übersetzen sich in formale Beweise im Mengenuniversum. ZFC wird zu einer Menge $\lceil \text{ZFC} \rceil$ von Aussagen auf der Objekzebene. $\text{Con}(\lceil \text{ZFC} \rceil)$ ist die Aussage, das es keinen Beweis in $\lceil \text{ZFC} \rceil$ von $\lceil \exists v_0 v_0 \neq v_0 \rceil$ gibt (hierbei steht „Con“ für engl. „consistent“).

Auf der Metaebene haben wir:

- (i) Ist ZFC widerspruchsfrei, so gilt $\text{non}(\text{ZFC} \vdash \text{CH})$.
- (ii) Ist ZFC widerspruchsfrei, so gilt $\text{non}(\text{ZFC} \vdash \neg \text{CH})$.

In die Objekzebene übersetzt sich das etwa als:

- (i) $\text{ZFC} \vdash \text{Con}(\lceil \text{ZFC} \rceil) \rightarrow \text{Con}(\lceil \text{ZFC} + \neg \text{CH} \rceil)$.
- (ii) $\text{ZFC} \vdash \text{Con}(\lceil \text{ZFC} \rceil) \rightarrow \text{Con}(\lceil \text{ZFC} + \text{CH} \rceil)$.

De facto läßt sich die Spiegelung der Metaebene in eine Objekzebene in einer viel schwächeren Theorie ausführen, und man kann ZFC links von $\vdash$ in (i) und (ii) etwa durch die Peano-Dedekind-Arithmetik PA ersetzen, wobei dann nicht nur Zeichen, sondern auch z. B. Formeln und Beweise als Zahlen kodiert werden müssen. Die Arithmetik kann sehr gut darüber nachdenken, was die Mengenlehre $\lceil \text{ZFC} \rceil$ beweisen kann.

Die Übersetzung von der Objekzebene in die Metaebene ist wesentlich problematischer. Hierzu betrachten wir den zweiten Gödelschen Unvollständigkeitssatz für ZFC. Er lautet schlicht und ergreifend:

- (+) Ist ZFC widerspruchsfrei, so gilt $\text{non}(\text{ZFC} \vdash \text{Con}(\lceil \text{ZFC} \rceil))$.

Im folgenden nehmen wir an, daß ZFC widerspruchsfrei ist. Dann ist nach (+) auch das schräge Axiomensystem $\text{ZFC}^* = \text{ZFC} + \neg \text{Con}(\lceil \text{ZFC} \rceil)$ widerspruchsfrei. Es gilt:

$\text{ZFC}^* \vdash \lceil \text{es gibt einen Beweis } b \text{ von } \lceil 0 = 1 \rceil \text{ in } \lceil \text{ZFC} \rceil \rceil$.

Ein Beweis b von $0 = 1$ in ZFC auf der Objekzebene kann aber nicht in einen Beweis von $0 = 1$ in ZFC auf der Metaebene übersetzt werden, d. h. es gibt keinen Beweis Δ mit $\lceil \Delta \rceil = b$. Andernfalls wäre ZFC widerspruchsvoll. Wie kann eine solche Übersetzung schiefgehen? Der Grund ist: b ist unendlich lang, die Objekzebene eines ZFC^* -Universums hat einen anderen Endlichkeitsbegriff als die Metaebene – und als das platonische Mengenuniversum V . „Das“ ω von ZFC^* , ω^* genannt, ist länger als ein platonisches ω , obwohl die Theorie ZFC^* felsenfest davon überzeugt ist, daß ω^* die erste Limesordinalzahl ist und genau die Menge der endlichen Ordinalzahlen darstellt.

Für den Platoniker ist $\text{Con}(\lceil \text{ZFC} \rceil)$ äquivalent mit der Widerspruchsfreiheit von ZFC. Zudem beweisen große Kardinalzahlen die Aussage $\text{Con}(\lceil \text{ZFC} \rceil)$. Axiomensysteme wie ZFC^* sind für den Platoniker lediglich ein weiteres Beispiel dafür, daß es wahre Aussagen gibt, die sich nicht in einer begrenzten Theorie wie ZFC beweisen lassen.

Man kann den zweiten Gödelschen Unvollständigkeitssatz in die Objektebene übersetzen und erhält dann:

(+) $\text{ZFC} \vdash \text{Con}(\ulcorner \text{ZFC} \urcorner) \rightarrow \text{Con}(\ulcorner \text{ZFC} + \neg \text{Con}(\ulcorner \text{ZFC} \urcorner) \urcorner)$.

Das ist interessant, aber unbedenklich. Sehr bedenklich wäre dagegen:

(-) ZFC ist widerspruchsfrei, aber $\text{ZFC} \vdash \neg \text{Con}(\ulcorner \text{ZFC} \urcorner)$.

Dieses metamathematische Kuriosum kann niemand als Nonsens entlarven, der ganz in ZFC verbleibt.

Der Platonist sagt: „Man lasse sich durch diese logischen Tricks im Umfeld der Gödel-resultate nicht verwirren, sondern eher in seiner platonischen Haltung bestärken. ZFC ist eine hübsche Sammlung von wahren Aussagen über das Universum V , die Sätze von ZFC erfassen aber nicht alle wahren Aussagen. Z. B. ist $\text{Con}(\ulcorner \text{ZFC} \urcorner)$ eine wahre Aussage, die nicht in ZFC beweisbar ist.“ Für den Erzplatonisten ist (CH) viel beunruhigender als ZFC^* -Spielereien: Er weiß nicht, ob (CH) in V gilt oder nicht. Er sucht nach mehr Wissen über Mengen, um irgendwann zu sehen, ob (CH) wahr ist oder falsch. ZFC gilt ihm als ein guter Anfang, aber sicher nicht als das letzte Wort.

Es gibt viele andere Dinge, die ZFC nicht beweisen kann, die ein Platoniker aber fast für selbstverständlich hält. Nach dem Gödelschen Vollständigkeitssatz gilt:

$\text{ZFC} \vdash \text{Con}(\ulcorner \text{ZFC} \urcorner) \leftrightarrow \exists \text{MM} \models \ulcorner \text{ZFC} \urcorner$.

Nach dem zweiten Unvollständigkeitssatz kann also ein widerspruchsfreies ZFC nicht zeigen, daß es ein Modell von $\ulcorner \text{ZFC} \urcorner$ gibt.

Der formalistische Standpunkt

Ein Gegenentwurf zum Platonismus ist die formalistische Haltung, die ihre Sicht der Mathematik aus der Möglichkeit eines finiten formalen Rahmens für die mathematische Tätigkeit gewinnt. Seine radikalste Form ist: Es gibt die mathematischen Objekte nicht wirklich, es gibt nur Zeichenketten auf dem Papier, und diese Zeichenketten werden nach festen Regeln manipuliert. Aus Axiomen – bestimmten ausgezeichneten Zeichenketten – gewinnt man durch diese Manipulationen mathematische Sätze.

Der Formalismus hat den Vorteil großer Klarheit. Er scheut die metaphysischen Grundannahmen des Platonismus wie der Teufel das Weihwasser, und gibt eine präzise Antwort auf die Frage, was Mathematik ist.

Die in diesem Buch eingenommene Haltung tendiert stark zum platonischen Entwurf. Die formale Welt ist uns ein willkommener, das Gebäude tragender Keller, der insbesondere metamathematische Aussagen ermöglicht.

Zur Illustration des formalistischen Standpunkts schließen wir dieses Kapitel mit einem neueren Kommentar zum Formalismus.

H. Dales über Formalismus

„... mathematicians are ambivalent between realism [Platonismus] and formalism. For example, I quote from Davis and Hersh (1981, p. 320):

... the typical working mathematician is a [realist] on weekdays and a formalist on Sundays. That is, when he is doing mathematics he is convinced that he is dealing with an objective reality whose properties he is attempting to determine. But then, when challenged to give a philosophical account of this reality, he finds it easiest to pretend that he does not believe in it after all.

... It seems to me that most mathematicians are formalists for all the days of the week. It is of course very useful when seeking proofs within the formal system to have a ‘realistic picture’ in one’s mind, and so it is temporarily convenient ... to be a realist ...

I think that the success of the major mathematicians in resolving problems and advancing the subject owes much to their ability to formulate in their mind an appropriate image of the abstract problem: it must be sufficiently subtle and complicated to capture the essential features of the question at issue, yet remain sufficiently simple to allow our limited minds absolutely and fully to explore, in quiet contemplation, all aspects of this image until we understand it sufficiently to begin the attempt to transfer this understanding to a written account of the general, abstract situation ...

Thus my view is that we are genuine, believing formalists who temporarily act as realists for reasons of expediency in solving problems.

... The first remark is that formalists practically never use a truly formal language in their writings (and may not know to do this, even under pressure); they formulate their theorems in the naive language of set theory developed in the XIXth century by Dedekind and Cantor. But they are confident that, if their results had to be formalized, this could be done; and doubtless they are correct in this.“

(H. Dales, „The mathematician as a formalist“. In: Dales, Olivieri 1998, „Truth in Mathematics“)

3. Mengen und Klassen

Dieses Kapitel bringt keine wesentlich neuen Gesichtspunkte der Axiomatik ZFC. Es gehört aber zur Diskussion des formalen Rahmens dieser Theorie, und zudem hat es aus ästhetischen Gründen hier seinen Platz, damit wir im zweiten Teil des Buches zur weiteren Entwicklung der axiomatischen Mengenlehre die bequeme Klassensprechweise bereits zur Verfügung haben und nicht mit einer Flut von Definitionen beginnen müssen. Am Ende des Kapitels besprechen wir kurz alternative Axiomensysteme, in denen echte Klassen echten Objektstatus genießen.

Außer Mengen gibt es in ZFC keine Objekte. Jedoch ist es bequem und suggestiv, auch über sogenannte Klassen reden zu können. Die Intuition ist hierbei: Mengen sind kleine Klassen, und Klassen, die keine Mengen sind – sogenannte echte Klassen – sind zu groß, um noch Mengen zu sein. Oder: Echte Klassen sind unbeschränkte Teile des Mengenuniversums – so wie unendliche Mengen von natürlichen Zahlen unbeschränkte Teile von $\mathbb{N}$ sind.

Bei der Diskussion der ZFC Axiome haben wir bewiesen:

$$\neg \exists z \forall y \ y \in z,$$
$$\neg \exists z \forall y. \ y \in z \leftrightarrow y \notin y.$$

Diese Ergebnisse können wir suggestiv schreiben als:

$$\neg \exists z \ z = \{ x \mid x = x \},$$
$$\neg \exists z \ z = \{ x \mid x \notin x \}.$$

Das Paarmengen- und das Potenzmengenaxiom besagen:

$$\forall x, y \exists z \forall u. \ u \in z \leftrightarrow u = x \vee u = y,$$
$$\forall x \exists z \forall u. \ u \in z \leftrightarrow u \subseteq x,$$

mit der Abkürzung $u \subseteq x = \forall z. \ z \in u \rightarrow z \in x$.

Auch diese Axiome können wir suggestiv notieren als:

$$\forall x, y \exists z \ z = \{ u \mid u = x \vee u = y \},$$
$$\forall x \exists z \ z = \{ u \mid u \subseteq x \}.$$

Allgemein können wir Ausdrücke der Form $\{ x \mid \varphi(x) \}$ – oder noch allgemeiner $\{ x \mid \varphi(x, p_1, \dots, p_n) \}$ für bestimmte Mengen (Parameter) $p_1, \dots, p_n$ – als Abkürzungen verwenden, etwa in den Formen

- (α) $y \in \{ x \mid \varphi(x) \},$
- (β) $z \subseteq \{ x \mid \varphi(x) \},$
- (γ) $\{ x \mid \varphi(x) \} \subseteq \{ x \mid \psi(x) \},$
- (δ) $z' = \{ x \mid \varphi(x) \},$

die nichts anderes bedeuten sollen als:

- (α) $\varphi(y)$,
- (β) $\forall x \in z \ \varphi(x)$,
- (γ) $\forall x. \varphi(x) \rightarrow \psi(x)$,
- (δ) $\forall x. x \in z' \leftrightarrow \varphi(x)$.

Klassen

Definition (*Klassen und die Konvention „Mengen sind Klassen“*)

- (i) Ist $\varphi(x, y_1, \dots, y_n)$ eine Formel der Mengenlehre, und sind $p_1, \dots, p_n$ Mengen, so heißt

$$\{ x \mid \varphi(x, p_1, \dots, p_n) \}$$

die zu φ gehörige Klasse (bzgl. $p_1, \dots, p_n$), gelesen: „die Klasse aller (Mengen) x mit $\varphi(x, p_1, \dots, p_n)$ “.

- (ii) Jede Menge y können wir wegen des Extensionalitätsaxioms als Klasse auffassen, indem wir y mit $\{ x \mid x \in y \}$ identifizieren.

Nach (ii) ist also jede Menge eine Klasse. Bei der Identifikation von y mit $\{ x \mid x \in y \}$ wird die Formel $\varphi(x, y) = x \in y$ mit y als Parameter verwendet.

Klassen bezeichnen wir in der Regel mit großen oder fettgedruckten Buchstaben $A, \mathbf{A}, B, \mathbf{B}, \dots$. Die Bezeichnung erfolgt oft in der Form $A = \{ x \mid \varphi(x) \}$.

Definition (*das mengentheoretische Universum V*)

$V = \{ x \mid x = x \}$ heißt die Allklasse oder das mengentheoretische Universum.

Wir verwenden Klassen informell als Objekte, streng genommen ist eine Klasse aber nichts anderes als eine Formel der Mengenlehre (evtl. mit Parametern). Klassen sind in diesem Sinn lediglich gut manipulierbare Sprachobjekte, und es wäre unbequem, auf diesen Komfort zu verzichten. Eine Aussage, die Klassen enthält, steht als Mitteilung für einen „echten“ Satz der Sprache $\mathcal{L}$ der Mengenlehre (vgl. die Klassen- und $\mathcal{L}$ -Formen von (α) – (δ) oben).

Definition (*$A = B$, $A \subseteq B$ für Klassen, echte Klassen*)

Seien $A = \{ x \mid \varphi(x) \}$, $B = \{ x \mid \psi(x) \}$ Klassen (aus Notationsgründen unterdrücken wir Parameter). Dann definieren wir:

- (i) $A = B$ als die Formel $\forall x. \varphi(x) \leftrightarrow \psi(x)$,
- (ii) $A \in B$ als die Formel $\exists x. x \in B \wedge x = A$,
- (iii) $A \subseteq B$ als die Formel $\forall x. x \in A \rightarrow x \in B$,
- (iv) A ist eine Menge als $A \in V$ (also als $\exists x \ x = A$),
- (v) A ist echte Klasse als $A \notin V$ (also als $\neg \exists x \ x = A$).

Sind A oder B Mengen, so stimmen diese Definitionen mit den alten überein. Insbesondere gilt dann für jede Menge x , daß $x = \{ y \mid y \in x \}$.

Weiter haben wir für Mengen x und Klassen $A = \{y \mid \varphi(y)\}$ die folgenden Auflösungen von Ausdrücken, die x und A enthalten:

$x \in A$ ist die Formel $\varphi(x)$,

$x = A$ ist die Formel $\forall z. z \in x \leftrightarrow z \in A$,

$A \in x$ ist die Formel $\exists z. z \in x \wedge z = A$,

$A \subseteq x$ ist die Formel $\forall y. y \in A \rightarrow y \in x$.

Schließlich sei wieder $\forall x \in A \varphi$ die Formel $\forall x. x \in A \rightarrow \varphi$ und $\exists x \in A \varphi$ die Formel $\exists x. x \in A \wedge \varphi$.

Die Antinomie von Russell und Zermelo wird, wie erwähnt, zum Satz der Mengenlehre: „ $R = \{x \mid x \notin x\}$ ist eine echte Klasse“. Ebenso ist V eine echte Klasse. Wir können die Antinomien der naiven Mengenlehre also interpretieren als: „Es gibt Mengen und echte Klassen.“

Übung

Seien A, B Klassen. Dann gilt:

- (i) $A \in B$ folgt $A \in V$ (d.h. eine Klasse hat nur Mengen als Elemente).
- (ii) Ist A eine echte Klasse und $A \subseteq B$, so ist B eine echte Klasse.

[Verwendung des Aussonderungsschemas.]

Operationen mit Klassen

Definition ($A \cap B, \bigcap A, A \times B$, usw. für Klassen)

Seien $A = \{x \mid \varphi(x)\}$, $B = \{x \mid \psi(x)\}$ Klassen (wieder ohne Angabe von Parametern). Dann definieren wir:

- (i) $A \cup B = \{x \mid x \in A \text{ oder } x \in B\} (= \{x \mid \varphi(x) \vee \psi(x)\})$,
- (ii) $A \cap B = \{x \mid x \in A \text{ und } x \in B\} (= \{x \mid \varphi(x) \wedge \psi(x)\})$,
- (iii) $A - B = \{x \mid x \in A \text{ und } x \notin B\} (= \{x \mid \varphi(x) \wedge \neg \psi(x)\})$,
- (iv) $\bigcup A = \{x \mid \text{es gibt ein } y \in A \text{ mit } x \in y\}$,
- (v) $\bigcap A = \{x \mid \text{für alle } y \in A \text{ gilt } x \in y\}$,
- (vi) $A \times B = \{x \mid \text{es gibt } a \in A \text{ und } b \in B \text{ mit } x = (a, b)\}$,
- (vii) $\mathcal{P}(A) = \{x \mid \text{für alle } y \in x \text{ gilt } y \in A\}$.

Ob man in $\{x \mid \varphi(x)\}$ die Formel φ umgangssprachlich notiert oder als $\mathcal{L}$ -Formel ist Geschmackssache. In jedem Falle wird man aber Abkürzungen verwenden wie „ $x = (a, b)$ “ oder „ x ist Funktion“, und damit um umgangssprachliche Formen nicht herumkommen. Dann ist etwa „ $\forall x \in y. x$ ist Funktion $\wedge \text{dom}(x) = \omega$ “ eine Mischform. Das gleichwertige „jedes $x \in y$ ist eine Funktion auf ω “ ist besser lesbar. Quantoren auszuschreiben ist oft schöner und hilft, syntaktische Wüsten und Privatnotationen zu vermeiden.

Für Mengen stimmen alle Definitionen mit den alten überein. Weiter gilt z.B.:

$$\bigcap \{A, B\} = A \cap B, \quad \bigcup \emptyset = \emptyset, \quad \bigcup V = V, \quad \bigcap V = \emptyset, \quad \bigcap \emptyset = V.$$

Relationen und Funktionen auf Klassen

Wir erweitern nun die Begriffe der Relationen und Funktionen auf Klassen.

Definition (*Klassen als Relationen und Funktionen*)

Seien A, B, C Klassen.

- (i) A heißt (*zweistellige*) *relationale Klasse auf B* oder kurz *Relation auf B* , falls $A \subseteq B \times B$, d. h. falls jedes $x \in A$ ein geordnetes Paar mit Elementen aus B ist.
- (ii) A heißt (*einstellige*) *funktionale Klasse* oder kurz *Funktion*, falls A Relation auf V ist und für alle x, y, z gilt:
 $(x, y) \in A$ und $(x, z) \in A$ folgt $y = z$.
- (iii) Eine Funktion A heißt *injektiv*, falls für alle x, y, z , gilt:
 $(x, z) \in A$ und $(y, z) \in A$ folgt $x = y$.
- (iv) Ist A Funktion, so heißt
 $\text{dom}(A) = \{ x \mid \text{es gibt ein } y \text{ mit } (x, y) \in A \}$ *der Definitionsbereich*
 [engl. domain] von A , und
 $\text{rng}(A) = \{ y \mid \text{es gibt ein } x \text{ mit } (x, y) \in A \}$ *der Wertebereich*
 [engl. range] von A .
 A heißt *Funktion auf B* , falls A Funktion und $\text{dom}(A) = B$.
- (v) A heißt *Funktion von B nach C* , in Zeichen $A : B \rightarrow C$, falls A Funktion auf B und $\text{rng}(A) \subseteq C$.
- (vi) $A : B \rightarrow C$ heißt *surjektiv*, falls $\text{rng}(A) = C$.
- (vii) $A : B \rightarrow C$ heißt *bijektiv*, falls A injektiv und surjektiv.

Ist F eine Funktion, so schreiben wir wie üblich $F(x) = y$ für $(x, y) \in F$. Analog sind mehrstellige Relationen A auf B ($A \subseteq B \times \dots \times B$) und mehrstellige Funktionen definiert, sowie Schreibweisen $F : A_1 \times A_2 \times \dots \times A_n \rightarrow B$, usw.

Ist F eine n -stellige Funktion, so schreiben wir im folgenden suggestiv

$\{ F(x_1, \dots, x_n) \mid \varphi(x_1, \dots, x_n) \}$ für

$\{ z \mid \text{es existieren } x_1, \dots, x_n \text{ mit } z = F(x_1, \dots, x_n) \text{ und } \varphi(x_1, \dots, x_n) \}$.

So ist etwa

$\{ (x, y) \mid x \in A \text{ und } y \in B \} =$

$\{ z \mid \text{es existieren } x \in A \text{ und } y \in B \text{ mit } z = (x, y) \} (= A \times B)$.

Auch die Elementrelation können wir als echte Klasse betrachten:

Definition

$\in = \{ (a, b) \mid a \in b \}$.

Häufig gebraucht werden:

Definition (*Bild, Einschränkung, Umkehrfunktion, Verkettung, Identität*)

Seien F, G Funktionen, und sei X eine Klasse.

- (i) $F''X = \{F(x) \mid x \in \text{dom}(F) \text{ und } x \in X\}$ heißt *das Bild von X (unter F)*.
- (ii) $F \upharpoonright X = F \cap (X \times V)$ heißt *die Einschränkung von F auf X* (genauer: auf $X \cap \text{dom}(F)$).
- (iii) Ist F injektiv, so heißt $F^{-1} = \{(y, x) \mid (x, y) \in F\}$ *die Umkehrfunktion von F* .
- (iv) $G \circ F = \{(x, z) \mid x \in \text{dom}(F), F(x) \in \text{dom}(G), z = G(F(x))\}$ heißt *die Verkettung von F und G* .
- (v) $\text{id} = \{(x, x) \mid x \in V\}$ heißt *die Identität*,
 $\text{id}_A = \text{id} \upharpoonright A$ *die Identität auf A* .

In der Literatur ist auch $F[X]$ für $F''X$ üblich.

Es gibt viele elementare Eigenschaften dieser Begriffe, z. B.:

- (α) Ist F injektiv, so ist $F^{-1} : \text{rng}(F) \rightarrow \text{dom}(F)$ bijektiv.
- (β) Sind $F : A \rightarrow B$ und $G : B \rightarrow C$ bijektiv,
so ist $G \circ F : A \rightarrow C$ bijektiv.
- (γ) Sind $F : A \rightarrow B, G : B \rightarrow C, H : C \rightarrow D$ Funktionen,
so ist $(H \circ G) \circ F = H \circ (G \circ F)$.
- (δ) Sind $F : A \rightarrow B$ und $G : C \rightarrow D$ Funktionen und ist
 $F \upharpoonright (A \cap C) = G \upharpoonright (A \cap C)$,
so ist $F \cup G : A \cup C \rightarrow B \cup D$ eine Funktion, usw.

Auch Relationen R und S kann man miteinander verketteten:

$R \circ_{\text{Rel}} S = \{(x, z) \mid \text{es gibt ein } y \text{ mit } (x, y) \in R \text{ und } (y, z) \in S\}$.

Sind F, G Funktionen, so gilt $G \circ F = F \circ_{\text{Rel}} G$.

Geordnete Paare

Eine kleine formale Schwierigkeit tritt auf bei der Behandlung von „Strukturen“ der Form $S = (A, B)$, wobei hier (zumeist) $B \subseteq A \times A$ gilt. Ist A eine echte Klasse, so können wir die übliche Paardefinition $(A, B) = \{\{A\}, \{A, B\}\}$ nicht verwenden, da Mengen nur Mengen als Elemente haben.

Wir verwenden aber dennoch *die Bezeichnung* (A, B) und sagen

„Sei (A, B) eine Struktur mit ...“ *anstelle des formal korrekten*

„Seien A, B Klassen, $B \subseteq A \times A$ und es gelte ...“, usw.

Alternativ kann man (A, B) definieren durch:

$$(A, B) = \begin{cases} \{\{A\}, \{A, B\}\}, & \text{falls } A \text{ und } B \text{ Mengen sind,} \\ A \times \{0\} \cup B \times \{1\}, & \text{falls } A \text{ echte Klasse oder } B \text{ echte Klasse.} \end{cases}$$

Häufig schreiben wir auch $\langle A, B \rangle$ für (A, B) .

Standardbeispiele sind $\langle A, \in \cap (A \times A) \rangle$, kurz $\langle A, \in \rangle$ oder auch $\langle A, \in \rangle$.

Aussagen mit Klassen und das Prinzip der Elimination

Ein Satz der Form „Sei A eine Klasse mit ... Dann gilt ...“ ist als ein Theoremschema anzusehen. So wie das Aussonderungsschema aus unendlich vielen Axiomen besteht – eines pro Formel –, besteht ein solches Theoremschema aus unendlich vielen Sätzen – je ein Satz pro Klasse/Formel. Analog ist ein Beweis eines Theoremschemas streng genommen ein Beweisschema. Ein Beweis von

„Zu einer Klasse A existiert genau eine Klasse B mit ...“

zeigt uns:

- (1) wie wir aus der A definierenden Formel ϕ eine B definierende Formel ψ mit den gewünschten Eigenschaften konstruieren können (effektiv auf dem Papier),
- (2) daß $B = B'$ gilt, falls B' eine weitere durch ψ' definierte Klasse mit den gewünschten Eigenschaften ist, d. h. der Beweis zeigt:
Für alle x gilt: $\psi(x) \text{ gdw } \psi'(x)$.

Setzt man in ein Theoremschema „Für alle Klassen A ...“ speziell $\phi = x \in y$ ein, d. h. bildet den Satz des Schemas für $A = \{x \mid x \in y\}$, so erhält man das Analogon des Schemas für Mengen. Die entsprechenden Sätze für Mengen werden also automatisch mitbewiesen.

Grundsätzlich gilt: Klassen lassen sich im Prinzip aus Definitionen und Theoremen immer entfernen, indem sie durch Formeln ersetzt werden. In der Praxis ist aber die Durchführung einer solchen Elimination nicht erforderlich und wäre wegen der suggestiven Stärke der Klassensprechweise auch nicht wünschenswert. Generell gilt: Keine Angst vor Klassen, auch nicht in ZFC.

Klassen in Rekursionen

Ein wichtiges Beispiel für eine Umwandlung von Formeln haben wir bereits im allgemeinen Rekursionssatz kennengelernt: Für jede Funktion $F : V \rightarrow V$ existiert genau eine Funktion G auf den Ordinalzahlen mit $G(\alpha) = F(G \upharpoonright W(\alpha))$ für alle Ordinalzahlen α . Ist $\phi(x, y)$ die definierende Formel von F , so können wir die definierende Formel $\psi(\alpha, x)$ von G aus dem Beweis des Rekursionssatzes ablesen. Sie lautet:

$\psi(\alpha, x) =$ „es gibt eine Funktion $g : W(\alpha + 1) \rightarrow V$ mit
 $\phi(g \upharpoonright W(\beta), g(\beta))$ für alle $\beta \leq \alpha$, und es gilt $x = g(\alpha)$ “.

ψ ist eine $\mathcal{L}$ -Formel. Der Beweis des Rekursionssatzes zeigt, daß diese Formel ψ eine Funktion G definiert, d. h. der Beweis zeigt: $\forall \alpha \exists! x \psi(\alpha, x)$, wobei $\forall \alpha =$ „für alle Ordinalzahlen α “. Weiter zeigt der Beweis, daß jede andere G definierende Formel zu ψ äquivalent ist.

In dieser Weise ist dann etwa $x \in V_\alpha$ eine $\mathcal{L}$ -Formel, mit der üblichen rekursiven Definition der V_α -Hierarchie. Wir müssen solche Formeln nicht ausschrei-

ben, aber im Prinzip könnten wir es tun. Wichtig ist, daß uns ZFC erlaubt, rekursive Definitionen über echte Klassen, etwa über alle Ordinalzahlen, durchzuführen. Dabei wird eine funktionale Klasse in eine andere übergeführt, und wir haben dann Zugriff auf die rekursiv definierten Objekte: Ganz so, wie wir für jedes α die Menge $\{\alpha\}$ betrachten können, können wir für jedes α die Menge V_α betrachten und damit arbeiten. Die resultierende Definition von V_α wollen wir als ein letztes konkretes Beispiel der Auflösung einer Rekursion noch einmal explizit ausschreiben:

$V_\alpha =$ „das eindeutige x mit:

es gibt eine Funktion $g : W(\alpha + 1) \rightarrow V$ mit:

$$g(0) = \emptyset,$$

$$g(\beta + 1) = \mathcal{P}(g(\beta)) \text{ für alle } \beta < \alpha,$$

$$g(\lambda) = \bigcup_{\beta < \lambda} g(\beta) \text{ für alle Limiten } \lambda \leq \alpha,$$

$$x = g(\alpha).“$$

Der Ausdruck: „ $z =$ das eindeutige x mit $\psi(x)$ “ meint zunächst, daß wir die Aussage $\exists! x \psi(x)$ beweisen können. In der weiteren Verwendung von z in Formeln ist dann „ $z = u$ “ oder „ $u = z$ “ identisch mit $\psi(u)$, „ $u \in z$ “ identisch mit $\exists x. \psi(x) \wedge u \in x$, und schließlich „ $z \in u$ “ identisch mit $\exists x. \psi(x) \wedge x \in u$. In dieser und keiner anderen Weise taucht dann etwa V_α in einer Formel auf, so etwa in: „für alle Ordinalzahlen α gilt $\alpha \subseteq V_\alpha$ “, nach ϵ -Auflösung der Abkürzung $\alpha \subseteq V_\alpha$. Nachdem man sich diese Dinge einmal klar gemacht hat, kann man dann mit diesem stillen Hintergrundwissen mit den in der Mengenlehre sehr häufig rekursiv definierten Objekten so frei umgehen wie mit allen anderen Objekten.

Klassen als echte Objekte

Der Leser mag sich fragen, warum Klassen nicht offiziell als Objekte zugelassen sind, wenn sie sich größtenteils vernünftig benehmen und man mit ihnen viele vertraute Operationen durchführen kann.

In der Tat liegt dies nicht daran, daß man Klassen nicht axiomatisch in den Griff bekommen würde: Es gibt alternative Systeme – NBG von Neumann-Bernays-Gödel und das System MK von Morse-Kelley –, die neben Mengen auch Klassen als Objekte zulassen. Wir werden diese Systeme unten kurz beschreiben.

ZFC hat aber neben seiner Natürlichkeit weitere Vorteile: Die Theorie ist in metamathematischen Untersuchungen einfacher zu handhaben. Zudem kann man Klassen zumeist ohne Scheu verwenden durch die „Klassen sind Formeln“-Konvention. Bei obigen Alternativen kann man sofort die (platonische) Frage nach Meta- oder Superklassen aufwerfen. Cantor hätte echte Klassen als Objekte sicherlich abgelehnt. Für ihn waren „inkonsistente Vielheiten“ zwar benennbar, aber keine echten Gegenstände der Theorie. Man kann über die „Vielheit aller Ordinalzahlen“ reden, und auch Operationen mit ihnen, etwa in „die Vielheit aller Paare von Ordinalzahlen“ machen Sinn, aber man hat die Ordinalzahlen Ω oder $\Omega \times \Omega$ nie als fertiges Ganzes. ZFC folgt in diesem Sinne der Cantorsche Sicht der Dinge.

Das System NBG

NBG steht für „von Neumann, Bernays, Gödel“. John von Neumann hat zwischen 1925 und 1928 eine Axiomatik entwickelt, in der echte Klassen als Objekte zugelassen sind. Paul Bernays untersuchte zwischen 1935 und 1954 ein ähnliches System, und Gödel formulierte 1940 eine auf den Ideen von von Neumann und Bernays ruhende Axiomatik mit Klassen. Diese Systeme lassen sich sehr natürlich in einer Weise formulieren, in der die Unterschiede zu ZFC minimalisiert sind. (Das originale System von von Neumann benutzt den Funktionsbegriff als Grundbegriff anstatt der Elementrelation.)

Die Sprache von NBG ist die von ZFC, also $\mathcal{L}$. Wir schreiben aber die Variablen nun in Großbuchstaben X, Y, Z , und nennen die Objekte der Theorie nun nicht Mengen, sondern Klassen. Wir definieren dann:

(+) X ist eine Menge *als die Formel* $\exists Y X \in Y$.

Eine Klasse X heißt *echte Klasse*, falls X keine Menge ist.

Für alle Formeln ϕ definieren wir dann die Mengenquantifikation durch:

$\forall x \phi$ *als die Formel* $\forall X. X \text{ ist eine Menge} \rightarrow \phi(X)$,

$\exists x \phi$ *als die Formel* $\exists X. X \text{ ist eine Menge} \wedge \phi(X)$.

Damit treten in Formeln Variablen x, y, z, X, Y, Z auf, die für Mengen bzw. Klassen stehen. Unsere Theorie enthält aber offiziell nur einen Typ von Objekten, genannt Klassen. Manche dieser Klassen sind Mengen, nämlich genau die, die sich in einer Klasse als Element finden. Wir haben definiert, was „ X ist Menge“ heißen soll, so wie wir in ZFC definiert haben, was $x \subseteq y$ bedeuten soll.

Die Axiome von NBG sind nun:

(EXT) Extensionalitätsaxiom

$$\forall X, Y. X = Y \leftrightarrow \forall z. z \in X \leftrightarrow z \in Y$$

(LM), (PA), (VER), (POT), (UN), (FUN)

Diese Axiome sind identisch mit denen von ZFC.

(KOM) Komprehensionsschema

$$\forall P_1, \dots, P_n \exists Y \forall u. u \in Y \leftrightarrow \phi(u, P_1, \dots, P_n)$$

wobei $\phi(u, P_1, \dots, P_n)$ eine Formel ist, die lediglich Mengenquantoren enthält, d. h. alle Quantoren in ϕ sind von der Form $\forall x$ oder $\exists x$.

(Dagegen erlauben wir allgemeine Klassenparameter $P_1, \dots, P_n$.)

(ERS) Ersetzungsschema

$$\forall F \forall x. F \text{ Funktion} \rightarrow \exists y \forall v. v \in y \leftrightarrow \exists u \in x (u, v) \in F.$$

wobei „ F Funktion“ = „ F ist rechtseindeutige Klasse von geordneten Paaren“ wie üblich definiert ist.

(AC) Auswahlaxiom

$$\exists F. F \text{ Funktion} \wedge \forall x. \exists y \in x \rightarrow \exists z \in x (x, z) \in F.$$

Der wesentliche Unterschied zu ZFC ist also: Wir dürfen volle Komprehension (über Formeln ohne Klassenquantifizierungen) bilden, und erhalten Objekte. Sind diese Objekte zu groß, so tauchen sie in keinen anderen Objekten mehr als Elemente auf. Die echten Klassen bilden in gewisser Weise den Abschluß der Mengewelt, man nimmt ihren Rand noch mit auf.

Es existiert dann etwa die Russell-Zermelo-Klasse $R = \{x \mid x \notin x\}$, liefert aber keinen Widerspruch mehr. Das übliche Argument zeigt lediglich, daß R eine echte Klasse ist. Es gilt $R \notin R$. Analog existiert die echte Klasse $V = \{x \mid x \text{ ist Menge}\} = \{x \mid \exists x \ x = x\}$ nach dem Komprehensionsschema.

Weiter haben wir in NBG eine starke Form des Auswahlaxioms, eine globale Auswahlfunktion F auf der Klasse aller nichtleeren Mengen $V - \{\emptyset\}$. Derartige globale Auswahlfunktionen F – genauer: Formeln, die F definieren – hat man in ZFC i. a. nicht (man hat sie aber z. B. in $ZFC + „V = L“$).

NBG beweist den folgenden Satz über Klassen, der im originalen System von von Neumann sogar ein Axiom ist, und der an Cantors Idee des „Hineinprojizierens“ erinnert (vgl. 2.13):

Satz (von-Neumann-Charakterisierung der Mengen in NBG)

Für alle X sind äquivalent:

- (i) X ist eine Menge.
- (ii) Es gibt kein $F : X \rightarrow V$ surjektiv (mit $V = \{x \mid x \text{ ist Menge}\}$).

NBG hat dagegen keine wirklich neuen Axiome. Für das System NBG gilt zudem, daß eine Aussage, die nur Mengenquantoren enthält, in NBG – (AC) genau dann beweisbar ist, wenn sie in ZF beweisbar ist [Shoenfield 1954], und das gleiche gilt für NBG und ZFC. Die Verwendung von echten Klassen in NBG und eine globale Auswahlfunktion liefern keine neuen Erkenntnisse über Mengen.

Andererseits ist NBG im Gegensatz zu ZFC endlich axiomatisierbar: NBG enthält zwar wie ZFC ein unendliches Axiomenschema in Form der Komprehension, jedoch gibt es ein zu NBG äquivalentes Axiomensystem bestehend aus Axiomen $A_1, \dots, A_n$. Damit ist dann $B = A_1 \wedge \dots \wedge A_n$ eine einzige zu NBG (und, was Mengen betrifft, zu ZFC) äquivalente Aussage, eine Art Weltformel.

Eine etwas andere Formulierung des Systems NBG verwendet eine Sprache $\mathcal{L}^*$ mit zwei Sorten von Variablen, nämlich Mengenvariablen $x, y, z, \dots$, und Klassenvariablen $X, Y, Z, \dots$. Axiome, die diese beiden Variablen miteinander in Verbindung bringen sind dann „jede Menge ist eine Klasse“ und „jede Klasse enthält nur Mengen als Elemente“, also $\forall x \exists Y \ x = Y$ und $\forall X, Y. X \in Y \rightarrow \exists z \ X = z$. Ansonsten ist das System analog.

Obwohl ZFC und NBG sehr ähnlich sind, ist an vielen Stellen Vorsicht geboten. So gilt z. B. in NBG nicht, daß für jede echte Klasse X jede Wohlordnung $\langle X, < \rangle$ von X (im Sinne der Diskussion in „Geordnete Paare“ oben) gleichlang mit $\langle \Omega, < \rangle$ ist mit

$\Omega = \{\alpha \mid \alpha \text{ ist eine Ordinalzahl (und eine Menge)}\}$.

$\langle X, < \rangle = \langle \Omega, < \rangle + \langle \Omega, < \rangle$ macht Sinn und ist länger als $\langle \Omega, < \rangle$.

Weiter gibt es in NBG i. a. keinen Beweis der φ -Induktion: Die Induktionsaussage $(\varphi(0) \wedge \forall n \in \omega. \varphi(n) \rightarrow \varphi(n+1)) \rightarrow \forall n \in \omega \varphi(n)$ ist i. a. in NBG nur für Formeln φ beweisbar, die lediglich Mengenquantoren enthalten [Mostowski 1951].

Das System MK

MK steht für „Morse, Kelley“. Es unterscheidet sich durch eine sehr naheliegende Liberalisierung der Komprehension. MK ist genau wie NBG, lediglich das Komprehensionsschema lautet nun:

(KOM) Komprehensionsschema

$\forall P_1, \dots, P_n \exists Y \forall u. u \in Y \leftrightarrow \varphi(u, P_1, \dots, P_n)$
wobei $\varphi(u, P_1, \dots, P_n)$ eine beliebige Formel ist.

Siehe hierzu [Morse 1965], den Anhang von [Kelley 1955] sowie [Wang 1949].

Das System MK ist freier als NBG, es gibt keine Verbote mehr bei der Komprehension, ein Formel-Check auf Klassenquantoren entfällt. Die Restriktion auf spezielle Formeln im Komprehensionsschema von NBG erscheint, öffnet man echten Klassen einmal die Tür, etwas halbherzig.

Das System MK ist im Gegensatz zu NBG nicht mehr endlich axiomatisierbar, und beweist neue Aussagen (metamathematischer Natur) über Mengen ([Kreisel / Levy 1968], [Kurata 1964], [Shepherdson 1951, 1952, 1953]). Das ist kein Grund, von ZFC auf eine Klassentheorie zu wechseln, denn die fehlenden Aussagen über Mengen lassen sich, ebenso wie die Konsistenz der Systeme NBG und MK, in Erweiterungen von ZFC um relativ schwache große Kardinalzahlaxiome beweisen (unerreichbare Kardinalzahlen oder Mahlo-Kardinalzahlen genügen). Erweiterungen von ZFC um solche große Kardinalzahlaxiome erlauben es, Modelle für mengentheoretische Klassentheorien wie NBG und MK in einer Objektwelt ohne Klassen zu untersuchen. Es gibt also keinen Grund für Neid und kein Gefühl des Mangels in einer liberal verstandenen Theorie, die keine Klassen kennt.

Mehr über NBG und MK findet der Leser neben den Originalarbeiten in [Fraenkel / Bar-Hillel / Levy 1973].

Vorläufiges Schlußwort

ZFC, angereichert um eine freie Klassensprechweise, ist bis heute das am meisten verwendete Axiomensystem für die Mengenlehre geblieben, und eignet sich besonders gut zur Gewinnung metamathematischer Resultate. Es ist flexibel, elegant, ausbaufähig, und genießt große Sympathien in weiten Kreisen der mathematischen Welt. Die Problemgeschichte der Mengenlehre spielte bei der Gestaltung von ZFC sicherlich eine große Rolle, allen voran Zermelos Beweis des Wohlordnungssatzes und die naheliegende „zu groß!“-Interpretation der Paradoxien der vollen Mengenkomprehension. Das iterative Mengenkonzept wird in ZFC interessanterweise sehr verschmitzt durch die Hintertür über das Ersetzungsschema und die V_α -Hierarchie implementiert, während sich das System nach außen für die Mathematik als Ganzes offen und natürlich gibt. ZFC ist kein Jahrgangschampagner für Kenner, der einen problemlosen Aufbau des Mengenuniversums begleitet, sondern ein Landwein aus recht bodenständigen Axiomen, der sich besonders auch bei formlosen Anlässen gut anbieten läßt. Aus mengen-

theoretischer Sicht ist zwischen dem Auswahlaxiom als Axiom und dem Wohlordnungssatz als Denkgesetz wenig Unterschied, aber psychologisch ist das Auswahlaxiom in seinen vielen „offensichtlichen“ Varianten ein rhetorisches Genie. Inwieweit ZFC als eine rustikale ad-hoc-Theorie oder eher als Shakespeare der Mathematik zu betrachten ist, und welche Rolle bühnenwirksamen Einzelaxiomen wie „ $V = L$ “ und dessen imposantem Gegenchor der großen Kardinalzahlen im mathematischen Theater des 21. Jahrhunderts zugeteilt werden wird, bleibt offen und aufregend.

Rückblickend bildet die Mengenlehre in der Handschrift von Georg Cantor für sich genommen bereits ein vollendetes Ganzes, eine Kathedrale der Mathematikgeschichte. Sie ist nichts, was man als Vorzeit des Eigentlichen abtun könnte. Ist die heutige axiomatische Mengenlehre mit ihrem formalen Fundament auch unstrittig eine würdige architektonische Weiterentwicklung der Cantorsche Ideen, so ist doch von einer detaillierten geschichtlichen Erforschung der für die moderne Mathematik so wesentlichen Jahrzehnte um 1900 noch viel Aufklärung und Anregung zu hoffen. Der Gang der Wissenschaft wie auch ihr zur Ruhe gekommener Bestandteil ist keineswegs so frei von Zufällen und menschlichen, kulturgeschichtlichen Motiven, wie es der Begriff der überprüfbaren Forschung zuweilen suggerieren möchte. Eine dominierende Theorie kann auch eine Blockade sein für etwas Neues innerhalb des Alten und außerhalb des Gewohnten. Cantor, der das Wesen der Mathematik nicht in ihrer Wahrheit oder Beständigkeit, sondern in ihrer Freiheit empfunden hat, trat für eine offene Mathematik ein, wenn er auch unablässig versucht hat, seine Kardinal- und Ordinalzahlen der Welt draußen ans Herz zu legen. Und auch dieser Text hofft, als Bühne für Cantor, Hausdorff und Zermelo, Zuneigung erweckt zu haben auch für eine geisteswissenschaftlich verstandene Mathematik und ihre Genese, ihre Einbettung in die Gezeiten des kulturellen Vor- und Zurückschreitens, für Menschen, die darum rangen, Gedanken für sich selber zu klären, um sie anderen vermitteln zu können: Alles ging, wie so oft, um Sprache. Alles mit dem Ziel, am Ende etwas auf der Zunge zu haben, was es zuvor nicht gab. Etwas, das höchsten Kriterien standhalten kann und nirgendwo hinzuschießen braucht um sich zu rechtfertigen. Die Unendlichkeit der Gedankenwelt ist Maß wie Feuerfunke für den schöpferisch tätigen Menschen. In der Mengenlehre wird diese Unendlichkeit selbst zum Thema, und die Mengenlehre ist letztendlich nur ein exemplarisches Studienmaterial für die, die die Kontemplation suchen des Seins und des Denkens.

*

*

*

4. Abschnitt

Anhänge



1. Liste der ZFC-Axiome

2. Lebensdaten der „dramatis personae“

3. Die wichtigsten Arbeiten von Cantor, Hausdorff und Zermelo

4. Zeittafel zur frühen Mengenlehre

5. Literatur

6. Notationen

7. Personenverzeichnis

8. Sachverzeichnis

1. Liste der ZFC-Axiome

Existenz der leeren Menge $\exists x \forall y y \notin x$

Extensionalitätsaxiom Zwei Mengen sind genau dann gleich, wenn sie die gleichen Elemente haben. $\forall x, y. x = y \leftrightarrow \forall z. z \in x \leftrightarrow z \in y$

Paarmengenaxiom Zu je zwei Mengen x, y existiert eine Menge z , die genau x und y als Elemente hat. $\forall x, y \exists z \forall u. u \in z \leftrightarrow u = x \vee u = y$

Vereinigungsmengenaxiom Zu jeder Menge x existiert eine Menge y , deren Elemente genau die Elemente der Elemente von x sind.
 $\forall x \exists y \forall z. z \in y \leftrightarrow \exists u. u \in x \wedge z \in u$

Potenzmengenaxiom Zu jeder Menge x existiert eine Menge y , die genau die Teilmengen von x als Elemente besitzt.
 $\forall x \exists y \forall z. z \in y \leftrightarrow \forall u. u \in z \rightarrow u \in x$

Aussonderungsschema Zu jeder Eigenschaft ϕ und jeder Menge x gibt es eine Menge y , die genau die Elemente von x enthält, auf die ϕ zutrifft.
 $\forall x, p_1, \dots, p_n \exists y \forall u. u \in y \leftrightarrow u \in x \wedge \phi(u, p_1, \dots, p_n)$

Ersetzungsschema Das Bild einer Menge unter einer Funktion ϕ ist eine Menge.
 $\forall p_1, \dots, p_n. \forall u \exists! v \phi(u, v, p_1, \dots, p_n) \rightarrow \forall x \exists y \forall v. v \in y \leftrightarrow \exists u. u \in x \wedge \phi(u, v, p_1, \dots, p_n)$

Unendlichkeitsaxiom Es existiert eine Menge x , die die leere Menge als Element enthält und die mit jedem ihrer Elemente y auch $\{y\}$ als Element enthält.
 $\exists x. (\exists y. y \in x \wedge \forall z z \notin y) \wedge \forall y. y \in x \rightarrow \exists z. z \in x \wedge \forall u. u \in z \leftrightarrow u = y$

Fundierungsaxiom Jede nichtleere Menge x hat ein Element y , das mit x kein Element gemeinsam hat. $\forall x. \exists y y \in x \rightarrow \exists y. y \in x \wedge \forall z. z \in y \rightarrow z \notin x$

Auswahlaxiom Ist x eine Menge, deren Elemente nichtleer und paarweise disjunkt sind, so existiert eine Menge y , die mit jedem Element von x genau ein Element gemeinsam hat.
 $\forall x. (\forall y (y \in x \rightarrow \exists z z \in y) \wedge \forall y, z. y \in x \wedge z \in x \wedge y \neq z \rightarrow \forall u. u \in y \rightarrow u \notin z) \rightarrow \exists y \forall z. z \in x \rightarrow \exists! u. u \in y \wedge u \in z$

Konvention: Verschieden bezeichnete Variablen sind verschieden.

2. Lebensdaten der „dramatis personae“

Georg Cantor	* 3.3.1845 Petersburg, † 6.1.1918 Halle
Felix Hausdorff	* 8.11.1868 Breslau, † 26.1.1942 Bonn
Ernst Zermelo	* 27.7.1871 Berlin, † 21.5.1953 Freiburg
Bernard Bolzano	* 5.10.1781 Prag, † 18.12.1848 Prag
Richard Dedekind	* 6.10.1831 Braunschweig, † 12.2.1916 Braunschweig
Julius König	* 16.12.1849 Raab (Győr), † 8.4.1913 Budapest
Arthur Schoenflies	* 17.4.1853 Landsberg/Warthe, † 27.5.1928 Frankfurt/M.
Giuseppe Peano	* 27.8.1858 bei Cuneo, † 20.4.1932 Turin
Ludwig Scheeffer	* 1.6.1859 Königsberg, † 11.6.1885 München
Ivar Bendixson	* 1.8.1861 Stockholm, † 29.11.1935 Stockholm (?)
Dimitry Mirimanov	* 13.9.1861 Perejaslavl Zalesski, † 5.1.1945 Genf
David Hilbert	* 23.1.1862 Königsberg, † 14.2.1943 Göttingen
Émile Borel	* 7.1.1871 Saint-Affrique, † 3.2.1956 Paris
Bertrand Russell	* 18.5.1872 Ravenscroft, † 2.2.1970 Penrhynedeudraeth
René Baire	* 21.1.1874 Paris, † 5.7.1932 Chambéry
Friedrich Hartogs	* 20.5.1874 Brüssel, † 18.8.1943 München
Gerhard Hessenberg	* 16.8.1874 Frankfurt/Main, † 16.11.1925 Berlin
Henri Lebesgue	* 28.6.1875 Beauvais, † 26.7.1941 Paris
Felix Bernstein	* 24.2.1878 Halle, † 3.12.1956 Zürich
Paul Mahlo	* 28.7.1883 Coswig, † 20.8.1971 Halle
Nikolai Lusin	* 9.12.1883 Irkutsk, † 25.2.1950 Moskau
Thoralf Skolem	* 23.5.1887 Sandsvaer, † 23.3.1963 Oslo
Abraham Fraenkel	* 17.2.1891 München, † 15.10.1965 Jerusalem
Mikhail Suslin	* 15.11.1894 Krasawka, † 1919 Moskau
Kazimierz Kuratowski	* 2.2.1896 Warschau, † 18.6.1980 Warschau
Pavel Alexandrov	* 7.5.1896 Bogorodsk, † 16.11.1982 Moskau
Alfred Tarski	* 14.1.1902 (1901 ?) Warschau, † 26.10.1983 Berkeley
John von Neumann	* 28.12.1903 Budapest, † 8.2.1957 Washington
Kurt Gödel	* 28.4.1906 Brünn, † 14.1.1978 Princeton
Stanisław Ulam	* 13.4.1909 Lemberg, † 13.5.1984 Santa Fe
Paul Cohen	* 2.4.1934 Long Branch (New Jersey)

3. Die wichtigsten Arbeiten von Cantor, Hausdorff und Zermelo

Cantor 1872 Über die Ausdehnung eines Satzes aus der Theorie der trigonometrischen Reihen

Cantor beweist eine schärfere Form seines Eindeutigkeitssatzes über die Entwicklung einer Funktion in eine trigonometrische Reihe aus dem Jahre 1870.

In der Arbeit findet sich auch die Cantorsche Definition der reellen Zahlen über Fundamentalfolgen (Cauchyfolgen) aus rationalen Zahlen. Er nimmt die rationalen Zahlen als Grundlage und definiert für Fundamentalfolgen (= reelle Zahlen) a, b die Beziehungen $a = b$, $a < b$, $a > b$ und die üblichen arithmetischen Operationen.

Cantor 1874 Über eine Eigenschaft des Inbegriffs aller reellen algebraischen Zahlen

Beweis der Abzählbarkeit der algebraischen Zahlen. Beweis der Überabzählbarkeit der reellen Zahlen durch Intervallschachtelung (siehe 1.8). Cantor betont, daß sich damit ein neuer Beweis der Existenz transzendenter Zahlen ergibt.

Cantor 1878 Ein Beitrag zur Mannigfaltigkeitslehre

Begriff der gleichen und kleineren Mächtigkeit beliebiger Mengen. Beweis der Gleichmächtigkeit verschiedendimensionaler Kontinua endlicher oder abzählbar unendlicher Dimension. Zur Konstruktion der Bijektion verwendet Cantor die Entwicklung irrationaler Zahlen in Kettenbrüche. Diskussion der Unstetigkeit der Bijektion und des Dimensionsbegriffs. Explizite Angabe der „Cantorsche Paarungsfunktion“ (vgl. 1.7). Erste Formulierung der Kontinuumshypothese.

Cantor 1879 Über unendliche, lineare Punktmannigfaltigkeiten (I)

Beginn der Untersuchung von Teilmengen reeller Zahlen (lineare Punktmengen). Begriffe: Isolierter Punkt einer Menge, Ableitung einer Punktmenge, „überall-dicht (in einem Intervall)“. Cantor iteriert die Operation der Ableitung und teilt die Punktmengen in zwei Klassen ein, je nachdem ob die Folge der Ableitungen nach endlich vielen Schritten in der leeren Menge terminiert oder nicht. Cantor wiederholt in dieser Arbeit zudem noch einmal den Mächtigkeitbegriff und den Beweis der Überabzählbarkeit der reellen Zahlen.

Cantor 1880 Über unendliche, lineare Punktmannigfaltigkeiten (II)

Die Arbeit markiert die Geburtsstunde der Ordinalzahlen. Cantor iteriert die Operation der Ableitung einer Punktmenge ins Transfinite. Er nennt die transfiniten Indizes der Iteration „Unendlichkeitssymbole“ und betrachtet z. B. ∞ , $\infty + 1$, $\infty + 2$, ..., 2∞ , $2\infty + 1$, $2\infty + 2$, ..., 3∞ , ..., ∞^2 , ..., ∞^3 , ..., ∞^∞ , ..., ∞^{∞^2} , ..., ∞^{∞^∞} , usw.

Cantor 1882 Über unendliche, lineare Punktmannigfaltigkeiten (III)

Cantor diskutiert seine extensionale Auffassung des Mengenbegriffs: Es genügt, wenn es „intern determiniert“ ist, ob ein Objekt einer Menge angehört oder nicht; es kommt nicht auf die „externe Determination“ an, d. h. darauf, ob man tatsächlich feststellen kann, welche der beiden Möglichkeiten zutrifft (etwa, ob eine gegebene reelle Zahl transzendent ist).

Er wiederholt den Abzählbarkeitsbegriff, und betont, daß die abzählbare Vereinigung abzählbarer Mengen wieder abzählbar ist.

Cantor zeigt, daß es im $\mathbb{R}^n$ höchstens abzählbar viele paarweise disjunkte stetige [zusammenhängende] Teilgebiete geben kann. Cantor benutzt zum Beweis nicht die Existenz einer abzählbar dichten Teilmenge – etwa $\mathbb{Q}^n$ –, sondern verwendet das Argument, mit dem heute üblicherweise gezeigt wird, daß es nicht überabzählbar viele paarweise disjunkte Mengen von positivem Maß geben kann.

In dieser Arbeit findet sich auch das Argument, daß beliebige Punkte im $\mathbb{R}^2$ durch „transzendente Pfade“ (sogar in Form von Kreisbögen) verbunden werden können, d. h. durch stetige Kurven, deren Punkte alle mindestens eine transzendente Koordinate haben (vgl. 1.8).

Cantor 1883 Über unendliche, lineare Punktmannigfaltigkeiten (IV)

Diese Arbeit bringt eine Fortsetzung der Untersuchung der Ableitungsoperation einer linearen Punktmenge. Weiter zeigt Cantor: Ist $[a, b]$ ein reelles Intervall und $P \subseteq [a, b]$ abzählbar, so läßt sich P durch endlich viele Intervalle mit beliebig kleiner Intervallsumme überdecken. (Das Argument ist relativ kompliziert; heute zeigt man den Satz sehr einfach mit Hilfe der Kompaktheit von $[a, b]$.)

Cantor 1883 Über unendliche, lineare Punktmannigfaltigkeiten (V)

Diese Arbeit ist die längste der Reihe und hat über weite Strecken philosophischen Charakter. Cantor führt hier die Fortsetzung der natürlichen Zahlenreihe ins Transfinite konsequent durch, und rechtfertigt den Begriff „Zahl“ für die neuen transfiniten Objekte.

Die transfiniten Zahlen teilt er in Zahlenklassen ein (§ 1), wobei die ersten beiden Zahlenklassen im Zentrum stehen: Die erste Klasse besteht aus den natürlichen (endlichen) Zahlen, die zweite aus den abzählbar unendlichen transfiniten Zahlen. Die Kontinuumshypothese lautet dann: Die reellen Zahlen haben die Mächtigkeit der zweiten Zahlenklasse.

Anschließend führt er den Begriff der „Wohlordnung“ (§§ 2 – 3) ein, und charakterisiert die transfiniten Zahlen als die Längen unendlicher Wohlordnungen. Er unterstreicht die fundamentale Bedeutung dieses Begriffes für die gesamte Mengenlehre. Die elementaren Sätze über Wohlordnungen sind ihm intuitiv klar (etwa der Vergleichbarkeitssatz: Von je zwei Wohlordnungen ist die eine ordnungsisomorph zu einem Anfangsstück der anderen). Er behandelt Addition und Multiplikation von Wohlordnungen.

In §§ 4 – 8 diskutiert Cantor die Begriffe „aktual unendlich“ und „potentiell unendlich“, und verteidigt das Konzept der aktualen Unendlichkeit und die Existenz der transfiniten Zahlen gegen verschiedene philosophische Einwände.

Weiter behandelt er noch einmal seine Konstruktion der reellen Zahlen von 1872 über Fundamentalfolgen und vergleicht seinen Ansatz mit den Konstruktionen von Dedekind und Weierstraß. Er sieht eine Notwendigkeit, das Kontinuum als „mathematisch-logischen Begriff“ und nicht als „religiöses Dogma“ zu behandeln (§§ 9 – 10). In einer Fußnote zu § 10 erwähnt er zum ersten Mal die heute nach ihm benannte Cantormenge C , das „Cantorsche Diskontinuum“.

In § 11 bespricht er Erzeugungsprinzipien für die transfiniten Zahlen (Nachfolgerbildung und Limesbildung). Die Zahlenklassen definiert er mit Hilfe eines Hemmungs- oder Beschränkungsprinzips, das ihm erlaubt, die Konstruktion der transfiniten Zahlen

an bestimmten Stellen zu stoppen, etwa dann, wenn zum ersten Mal eine überabzählbare Zahl erreicht wird. In § 12 zeigt er, daß die zweite Zahlklasse eine größere Mächtigkeit hat als die erste, in § 13, daß diese Mächtigkeit minimal größer als die Mächtigkeit der ersten Zahlklasse ist. § 14 schließlich beschäftigt sich mit der Arithmetik der transfiniten Zahlen.

Cantor 1884 Über unendliche, lineare Punktmannigfaltigkeiten (VI)

Diese Arbeit ist die direkte Fortsetzung von (V), und enthält fünf weitere Paragraphen.

In §§ 15 – 18 wird die transfinit-iterierte Ableitung von Punktmengen behandelt. Cantor definiert den Begriff der Abgeschlossenheit: $P \subseteq \mathbb{R}$ heißt abgeschlossen, falls die Ableitung von P [Menge ihrer Häufungspunkte] eine Teilmenge von P ist. Wie schon in Teil (V) heißt eine Menge perfekt, falls sie mit ihrer Ableitung übereinstimmt (also perfekt = abgeschlossen + keine isolierten Punkte). Das Hauptergebnis dieser Thematik ist der Satz von Cantor-Bendixson: Ist $P \subseteq \mathbb{R}$ überabzählbar und abgeschlossen, so existiert eine eindeutige Darstellung $P = Q \cup A$ mit $Q \cap A = \emptyset$, Q perfekt, A abzählbar.

§ 18 leitet die Maßtheorie ein: Cantor definiert erstmalig in der Geschichte einen Inhalt für beliebige Teilmengen im $\mathbb{R}^n$. Der hier vorgestellte Ansatz wurde von Jordan und Peano weiterverfolgt und führte zum endlichen additiven Peano-Jordan-Inhalt. Borel und Lebesgue entwickelten die σ -additive Maßtheorie.

§ 19 beschäftigt sich wieder mit perfekten Teilmengen von $\mathbb{R}$. Cantor zeigt, daß jede perfekte Menge die Mächtigkeit von $\mathbb{R}$ hat. Hierzu verwendet er die Cantormenge [1883b], und zeigt, daß jede beschränkte nichtleere perfekte Teilmenge von $\mathbb{R}$, die keine nichttrivialen Intervalle enthält, ähnlich ist zur Cantormenge C . Zusammen mit dem Satz von Cantor-Bendixson ist damit die Kontinuumshypothese für die abgeschlossenen Mengen bewiesen: Jede abgeschlossene Menge ist endlich, abzählbar unendlich oder von der Mächtigkeit von $\mathbb{R}$.

Cantor 1892 Über eine elementare Frage der Mannigfaltigkeitslehre

Beweis von $|M| < |\mathcal{P}(M)|$ für alle Mengen M mit Hilfe des Diagonalverfahrens: Cantor beweist, daß die Menge aller Funktionen von M nach $\{0, 1\}$ von größerer Mächtigkeit ist als die Menge M .

Cantor 1895 Beiträge zur Begründung der transfiniten Mengenlehre (I)

Die zwei Teile der „Beiträge“ bilden Cantors abschließende, systematische Darstellung seiner Mengenlehre. Diese etwa 70 Seiten sind das erste Lehrbuch der Mengenlehre.

Die Arbeit beginnt mit seiner berühmten Mengendefinition der „Zusammenfassung zu einem Ganzen“. In §§ 1 – 6 wird dann der Mächtigkeitsbegriff behandelt (Größenvergleiche, Arithmetik mit Mächtigkeiten, Endlichkeit, die kleinste unendliche Mächtigkeit $\aleph_0$). § 7 führt den Begriff der einfach geordneten Mengen [totale oder lineare Ordnungen] ein, § 8 behandelt ihre Addition und Multiplikation. § 9 bringt ein neues Resultat über die Charakterisierung des Ordnungstyps der rationalen Zahlen: „Jede dichte, unbeschränkte und abzählbare lineare Ordnung ist ordnungsisomorph zu den rationalen Zahlen“. § 10 untersucht unendliche aufsteigende und absteigende Folgen in linearen Ordnungen (Fundamentalreihen). In § 11 charakterisiert Cantor die reellen Zahlen rein ordnungstheoretisch. (Die heutige Form ist: Jede vollständige, unbeschränkte lineare Ordnung, die eine abzählbare dichte Teilmenge besitzt, ist ordnungsisomorph zu den reellen Zahlen; Cantor verwendet „perfekt“ statt „vollständig“.)

Cantor 1897 Beiträge zur Begründung der transfiniten Mengenlehre (II)

§§ 12 – 14 beschäftigen sich mit Wohlordnungen und Ordnungszahlen [Ordinalzahlen, oder, im unendlichen Fall, transfiniten Zahlen]. § 15 führt „die zweite Zahlklasse“ der ab-

zählbaren Ordnungszahlen ein, und in § 16 wird die Gleichmächtigkeit der zweiten Zahlklasse mit der „zweitkleinsten Kardinalzahl $\aleph_1$ “ gezeigt. In § 17 zeigt Cantor die Cantorsche Normalform der Darstellung von abzählbaren Ordnungszahlen. § 18 beschäftigt sich mit der Exponentiation von Ordnungszahlen. Die Exponentiation wird durch transfinite Rekursion definiert. In § 19 wird ein allgemeinerer Satz über Normalformen der Darstellung abzählbarer Ordnungszahlen bewiesen. § 20 schließlich behandelt die sogenannten ϵ -Zahlen (Fixpunkte der Ordinalzahlexponentiation $f(\alpha) = \omega^\alpha$).

Hausdorff 1904 Der Potenzbegriff in der Mengenlehre

Die erste mengentheoretische Arbeit von Hausdorff, dreiseitig und skizzenhaft, mit „aus dem Sprechsaal“ überschrieben. Hausdorff gibt die „Hausdorff-Formel“ der Kardinalzahlarithmetik an. Interessant ist weiter die Behauptung der Regularität von Nachfolgerkardinalzahlen, aus der sich der Multiplikationssatz für Wohlordnungen ergeben würde. Bewiesen wurde dieser dann aber vollständig wohl erst von Hessenberg 1906.

Hausdorff 1906 – 1907 Untersuchungen über Ordnungstypen I – V

Hausdorff 1907 Über dichte Ordnungstypen

Hausdorff 1908 Grundzüge einer Theorie der geordneten Mengen

In diesen Arbeiten untersucht Hausdorff lineare Ordnungen, ihre Typen und ihre Arithmetik. Ziel war eine Klassifikation dieser Typen, wie sie Cantor vollständig für die Wohlordnungen gelang. Hausdorff bewies hier wichtige Struktursätze, etwa über dichte und zerstreute Ordnungen und Lücken in linearen Ordnungen. Daneben entwickelte er eine allgemeine Potenzierungstheorie, die z. B. zur Hausdorff-Hessenberg-Darstellung der Ordinalzahlpotenz führte.

Die Vielzahl von Begriffen und Ergebnissen über lineare Ordnungen kann hier nicht wiedergegeben werden. Einiges findet der Leser unten zu Kapitel 4 und 6 der „Grundzüge der Mengenlehre“, insbesondere zu Verallgemeinerungen des Typs $\eta = \eta_0$ der rationalen Zahlen. Diese sogenannten η_α -Typen sind intuitiv an jeder Stelle besonders dicht gepackte lineare Ordnungen.

Bemerkenswerte Eigenarten der Arbeiten sind ihr von den umhergeisternden Paradoxien unbeeinträchtigt Stil und ihre schlafwandlerisch sichere Weiterentwicklung der Cantorschen Begriffe und Ideen. Hausdorff hat mit einem Schiff der Cantorschen Mengenlehre neue Kontinente entdeckt, während fast alle anderen die Flotte vor dem Sinken retten wollten. Sogar Paul Mongré mußte vor Kapitän Hausdorff zurücktreten. In seinem Buch von 1914 hat dieser dann die meisten seiner Entdeckungen kartographiert.

Hausdorff 1914 Grundzüge der Mengenlehre

Das Buch hat zehn Kapitel, und läßt sich in zwei größere Themengebiete unterteilen: Allgemeine Mengenlehre mit einem Schwerpunkt auf der Ordnungstheorie (Kapitel 1 – 6) und daneben die mathematische Theorie des Raumbegriffs: topologische und metrische Räume, stetige Funktionen und Funktionenfolgen, Maßtheorie (Kapitel 7 – 10).

Hausdorff 1914 Grundzüge der Mengenlehre, Kapitel 1:

Mengen und ihre Verknüpfungen: Summe, Durchschnitt, Differenz.

Das Kapitel gibt eine Einführung in den „naiven Mengenbegriff“. Hausdorff bespricht den iterativen Charakter der Mengen: Mengen lassen sich zu neuen Mengen zusammenfassen. Weiter betont er die für die Mathematik fundamentstiftende Kraft des Konzepts der unendlichen Menge. Auf die noch „nicht abgeschlossenen“, „äußerst scharfsinnigen

Untersuchungen“ von Ernst Zermelo will er aber nicht weiter eingehen. Das Kapitel geht dann also ohne die axiomatische Bewaffnung die ersten Schritte der Mengenlehre: Schnitt, Vereinigung, Differenzen, Mengensysteme (etwa Ringe, σ -Systeme), Dualität, limes inferior und limes superior von Mengenfolgen, Folgen reeller Zahlen, usw.

Hausdorff 1914 Grundzüge der Mengenlehre, Kapitel 2:
Mengen und ihre Verknüpfungen: Funktion, Produkt, Potenz.

Kapitel 2 bringt die erste rein mengentheoretische Definition von „Funktion“, basierend auf einer Definition des geordneten Paares (a, b) als Menge $\{\{a, 1\}, \{b, 2\}\}$. (Diese wurde 1921 von Kuratowski zum heute üblichen geordneten Paar $(a, b) = \{\{a\}, \{a, b\}\}$ verbessert.) Nach einer Definition der Abzählbarkeit folgen die Betrachtung von „großen“ Schnitten, Vereinigungen und Produkten, und eine Diskussion von Rechenregeln, etwa der Distributivgesetze.

Hausdorff 1914 Grundzüge der Mengenlehre, Kapitel 3:
Kardinalzahlen oder Mächtigkeiten.

Hausdorff hat sich für die Durchführung all der hübschen Motive der ersten beiden Kapitel immerhin 45 reich instrumentierte Seiten gegönnt. Nun folgt ein Kapitel über Mengenlehre im eigentlichen Sinn. Hausdorff referiert die Cantorsche Auffassung von Kardinalzahlen als Ergebnis eines Abstraktionsprozesses, kritisiert die Russellsche Definition, die echte Klasse aller zu einer Menge gleichmächtigen Mengen als Kardinalzahl zu definieren, und gibt sich dann selber mit dem „formalen Standpunkt“ zufrieden, bei dem jeder Menge ein Ding zugeordnet wird in der Weise, daß gleichmächtigen Mengen das gleiche Ding entspricht. Und genau diese „Dinge“ werden dann per Ritterschlag als Kardinalzahlen oder Mächtigkeiten geadelt. Das ist gerade formal problematisch, aber Hausdorff kann damit überaus bequem und fehlerfrei arbeiten. Auf diese sehr abstrakten Zuordnungen folgt zur Erholung von Autor und Leser die Erläuterung des Vergleichs von Mächtigkeiten durch eine Paarbildung von „Äpfeln und Birnen“, ähnlich unseren Nußhaufen in 1.4. Innerhalb der Diskussion des Vergleichbarkeitsproblems beweist Hausdorff den Satz von Cantor-Bernstein-Dedekind; er gibt einen Beweis mit natürlichen Zahlen nach Bernstein und einen zahlenfreien nach Dedekind-Zermelo. Danach werden Summe, Produkt und Potenz von Mächtigkeiten definiert. Hausdorff beweist nun: $|\mathbb{N}| \leq |M|$ für jede unendliche Menge M , den Satz von Cantor über die Potenzmenge, und den Satz von König-Zermelo über Summe und Produkt von Kardinalzahlen. Schließlich werden die Kardinalitäten $\aleph_0, 2^{\aleph_0}, 2^{2^{\aleph_0}}$, also die Mächtigkeiten von $\mathbb{N}, \mathbb{R}$ und $\mathfrak{F}$ ins Auge gefaßt und die üblichen Gleichungen bewiesen, z. B. etwa $2^{\aleph_0} \cdot 2^{\aleph_0} = 2^{\aleph_0}$. Hausdorff notiert das Kontinuumsproblem, und bemerkt, daß es „bis jetzt allen Bemühungen widerstanden“ hat.

Hausdorff 1914 Grundzüge der Mengenlehre, Kapitel 4:
Geordnete Mengen. Ordnungstypen.

Start der Untersuchung von geordneten Mengen, die sich über drei Kapitel erstreckt. Lineare Ordnungen werden als Mengen von geordneten Paaren eingeführt, die (in heutiger Sprache) die Bedingungen der Linearität, Irreflexivität und Transitivität erfüllen. Die Definition ist so ein unmittelbarer Vorläufer der heutigen Definition einer Ordnung als Struktur $\langle M, < \rangle$. Hausdorff definiert nun die Ähnlichkeit (Ordnungsisomorphie) zweier linearer Ordnungen, und führt anschließend analog zu seinem Kardinalzahlbegriff den Begriff des Ordnungstypus ein als das allen ordnungsisomorphen Mengen Gemeinsame: Wieder denkt er sich jeder geordneten Menge ein „Zeichen“ α an die Seite gestellt derart, daß isomorphe Ordnungen das gleiche Zeichen erhalten. Anschließend werden Verknüp-

funken von Ordnungen betrachtet, etwa die Summenbildung und Produkte mit endlich vielen Faktoren. Es folgt eine Diskussion von Strecken und Stücken einer geordneten Menge. Danach wendet sich Hausdorff der Feinstruktur einer linearen Ordnung zu, und bespricht Begriffe wie „dicht“, „Lücke“, „Stetigkeit“, „zerstreut“ usw. Weiter wird hier auch der heute zentrale Begriff der Konfinalität besprochen, der 1906 von Hausdorff und Hessenberg eingeführt wurde: Eine Teilmenge N einer linear geordneten Menge M heißt konfinal in M , falls für jedes $x \in M$ ein $y \in N$ existiert mit $x \leq y$. Die Bezeichnung *konfinal* stammt von Hausdorff. Die Konfinalität von M ist dann die kleinste Kardinalität einer konfinalen Teilmenge von M . Sie gibt die minimale Anzahl der Schritte an, die benötigt werden, um in einer linearen Ordnung „von links nach rechts gehend“ bis an ihr Ende zu gelangen.

Im letzten Abschnitt des Kapitels wird gezeigt, daß es genau $\mathbb{R}$ -viele Ordnungstypen von abzählbaren linearen Ordnungen gibt. Diesen Satz hatten Bernstein und Hausdorff aufbauend auf Anregungen von Cantor unabhängig voneinander bereits 1901 gefunden.

Hausdorff 1914 Grundzüge der Mengenlehre, Kapitel 5:
Wohlgeordnete Mengen. Ordnungszahlen.

Hausdorff definiert Wohlordnungen als lineare Ordnungen, deren „jedes von Null verschiedene Endstück ein erstes Element hat“. Die Ordnungszahlen sind dann die Ordnungstypen von Wohlordnungen. Es folgt ein Beweis des Vergleichbarkeitssatzes. Hausdorff führt hier die „durchweg festgehaltene“ Notation „ $W(\alpha)$ = Menge der Ordinalzahlen $< \alpha$ “ ein, die bei der Neumann-Zermelo-Definition überflüssig wird, unter der $W(\alpha) = \alpha$ gilt. Hausdorff gibt eine arithmetische Liste für die ersten Ordinalzahlen nach Cantorschem Vorbild. Es folgt ein Abschnitt über transfinite Induktion, in welchem auch Normalfunktionen, d. h. monotone und stetige Funktionen auf den Ordnungszahlen untersucht werden. Hausdorff definiert ihre Fixpunkte als „kritische Zahlen“. Die Untersuchung des Phänomens geht auf Cantor (1897), Hessenberg (1906) und Oswald Veblen (1908) zurück.

Weiter geht es mit Ordinalzahlarithmetik. Danach werden Alephs als die Mächtigkeiten unendlicher Wohlordnungen eingeführt, und die mit Ordnungszahlen indizierte Folge aller Alephs betrachtet. Hausdorff beweist hierauf den Hessenbergschen Multiplikationssatz von 1906 in einer Variante des Jourdainischen Beweises von 1908.

Hausdorff definiert weiter reguläre Ordnungszahlen: α ist regulär, falls die Konfinalität von α gleich α ist; die Konfinalität ist hierbei die kleinste Mächtigkeit einer in $W(\alpha)$ konfinalen Teilmenge (vgl. Kapitel 4 oben). Hausdorff zeigt, daß alle Alephs der Form $\aleph_{\alpha+1}$ regulär sind. Er stellt wie schon 1908 die Frage nach der Existenz regulärer Alephs mit Limesindex, und schreibt:

„Wenn es also reguläre [Alephs] mit Limesindex gibt (und es ist bisher nicht gelungen, in dieser Annahme einen Widerspruch zu entdecken), so ist die kleinste von ihnen von einer so exorbitanten Größe, daß sie für die üblichen Zwecke der Mengenlehre kaum jemals in Betracht kommen wird.“

Die Rede ist von schwach unerreichbaren Kardinalzahlen, die heute die kleinsten großen Kardinalzahlen sind – „groß“ ist immer relativ –, und die, wie ihre viel größeren Verwandten, für die Zwecke der Mengenlehre von zentraler Bedeutung sind. Ihre Existenz ist in ZFC nicht beweisbar, sie sind aller Wahrscheinlichkeit nach widerspruchsfrei, und führen mittlerweile eine zumindest platonisch gutgesicherte Existenz. Der Name „unerreichbar“ geht auf Kuratowski zurück, und zeigt wie der Ausdruck „exorbitant“, wie viel Respekt man dem David unter den großen Kardinalzahlen damals entgegenbrachte.

Am Ende des sehr reichhaltigen Kapitels wird als Desert dann auch noch der Wohlordnungssatz von Zermelo bewiesen.

Hausdorff 1914 Grundzüge der Mengenlehre, Kapitel 6: Beziehungen zwischen geordneten und wohlgeordneten Mengen.

Das siebzig Seiten starke Kapitel beginnt mit einem Ordnungsbegriff, der heute in der Mathematik überall zu finden ist: Hausdorff definiert teilweise geordnete Mengen (partielle Ordnungen). Hausdorff beweist in diesem Kapitel das heute nach ihm benannte Maximalitätsprinzip, aus dem sich z. B. das „Zornsche Lemma“ unmittelbar gewinnen läßt.

Hausdorff definiert das Produkt $A = \mathfrak{P}_{i \in J} A_i = \{ f \mid f: J \rightarrow \bigcup_{i \in J} A_i, f(i) \in A_i \text{ für } i \in J \}$ von linearen Ordnungen über einer linear geordneten Indexmenge J . Die Elemente dieses Produkts können durch Vergleich an der Stelle ihres kleinsten Unterschiedes geordnet werden, falls eine solche kleinste Stelle existiert. Es entsteht so eine partielle Ordnung (lexikographische Ordnung von A). Ist J wohlgeordnet, so ist die Ordnung auf dem Produkt linear.

Im allgemeinen Fall linearer Indexmengen J kann Hausdorff eine maximale lineare Teilmenge der partiellen Ordnung auf dem Produkt *definieren* (ohne Rückgriff auf sein Maximalitätsprinzip).

Weiter untersucht Hausdorff Verallgemeinerungen des Ordnungstyps $\eta = \eta_0$ der rationalen Zahlen. Für Leser mit speziellem Interesse hierzu zwei Begriffe. Zunächst definiert Hausdorff η_α (für ein $\alpha \geq 0$) als den Typus von $M_\alpha = \{ f \mid f: W(\omega_\alpha) \rightarrow \{0, 1, 2\}, f \text{ ist schließlich konstant gleich } 1 \}$, lexikographisch geordnet. Weiter definiert Hausdorff: Eine lineare Ordnung heißt η_α -Menge, falls zwischen einer monoton aufsteigenden Folge der Länge $< \aleph_\alpha$ und einer darüberliegenden monoton absteigenden Folge der Länge $< \aleph_\alpha$ immer noch Punkte der Ordnung liegen. Solche Folgen laufen also niemals auf eine Lücke der Ordnung zu. Offenbar sind die rationalen Zahlen η_0 -Mengen, und ihr Typ ist η_0 . Hausdorff zeigt insbesondere, daß für jedes α gilt:

1. Jede Menge vom Typ $\eta_{\alpha+1}$ ist eine η_α -Menge der Größe $2^{\aleph_\alpha}$.
2. Je zwei η_α -Mengen der Mächtigkeit $\aleph_\alpha$ sind ordnungsisomorph.
3. Jede Ordnung der Größe $\aleph_\alpha$ läßt sich in jede η_α -Menge ordnungstreu einbetten.

Die η_α -Mengen haben also eine bemerkenswerte Universalitätseigenschaft. Der Leser vergleiche die Aussagen mit der Charakterisierung des Ordnungstyps $\eta = \eta_0$ durch Cantor. Für die Existenz von $\eta_{\alpha+1}$ -Mengen der Mächtigkeit $\aleph_{\alpha+1}$ ist $2^{\aleph_\alpha} = \aleph_{\alpha+1}$ notwendig und hinreichend. Gilt (GCH), so bergen die sehr einfach definierten lexikographischen Ordnungen $M_{\alpha+1}$ der Größe $\aleph_{\alpha+1}$ *alle* linearen Ordnungen der Größe $\aleph_{\alpha+1}$ in sich, ganz so, wie sich in die rationalen Zahlen alle abzählbaren Ordnungen hineinkopieren lassen.

Das Kapitel widmet sich dann noch „rationalen Ordnungszahlen“. Danach werden Alternativen zur lexikographischen Ordnung auf dem Produkt von linearen Ordnungen diskutiert, und der Spezialfall eines Produktes betrachtet, dessen Faktoren alle die linearen Ordnungen der reellen Zahlen sind.

Hausdorff 1914 Grundzüge der Mengenlehre, Kapitel 7: Punktmengen in allgemeinen Räumen.

Der zweite Teil des Buches beginnt mit diesem Kapitel, und das Kapitel selbst mit den Worten:

„In der Anwendung auf die Punktmengen des Raumes, in der Klärung und Verschärfung der geometrischen Grundbegriffe hat die Mengenlehre ihre schönsten Triumphe gefeiert, die selbst von denjenigen zugestanden werden, die sich der abstrakten Mengenlehre gegenüber skeptisch verhalten.“

Hausdorff motiviert mit diesem Loblied zu Beginn eines Themenwechsels sowohl die von der abstrakten Ordnungstheorie erschöpften Leser, als auch diejenigen, die gerne mehr von der abstrakten Mengenlehre gehört hätten. Und in der Tat beginnt nun etwas

ganz anderes, auch wenn Hausdorff sich Mühe gibt, den Übergang etwas weicher zu machen. Nach wenigen Absätzen, die den Wunsch nach möglichst allgemeinen Definitionen der Raumbegriffe begründen, definiert Hausdorff „metrische Räume“ über „Entfernungssaxiome“ und danach „topologische Räume“ über „Umgebungsaxiome“. Mit der Wortwahl „topologisch“ will Hausdorff „andeuten, daß es sich um Dinge handelt, die ohne Maß und Zahl ausdrückbar sind“. Während den metrischen Räumen eine reellwertige Abstandsfunktion zwischen den Punkten zugrunde liegt, sind es bei topologischen Räumen nur die Beziehungen, die abstrakte „Umgebungen“ eines Punktes untereinander erfüllen. Hausdorffs Definition eines metrischen Raumes ist die heute übliche, seine Definition eines topologischen Raumes hat sich später im wesentlichen durch eine Bearbeitung durch Heinrich Tietze (31. 8. 1880 – 17. 2. 1964) 1923 zur heutigen Definition vollendet. Bemerkenswert ist, daß Hausdorff, der in seinem Buch wie anderswo die Zermelo-Axiomatik für die Mengenlehre nicht diskutiert, seine Raumlehre in eine axiomatische Form gießt. Die intendierte Vielzahl von Modellen für metrische und topologische Räume im Gegensatz zur eindeutig empfundenen Welt der Mengen mag hierfür der Grund sein.

Im folgenden diskutiert Hausdorff innere Punkte, Randpunkte, „Gebiete“ (heute: offene Mengen; Hausdorffs topologische Räume sind über Umgebungsbasen definiert, noch nicht über offene Mengen), Berührungspunkte, Häufungspunkte, Verdichtungspunkte, die Cantorsche Begriffe „abgeschlossen, insichdicht und perfekt“, kompakte Mengen nach Fréchet (Mengen, bei denen jede unendliche Teilmenge einen Häufungspunkt hat, „abgeschlossen“ ist hier noch nicht dabei) und beweist zugehörige elementare Sachverhalte, insbesondere die Überdeckungseigenschaft (Satz von Borel), die heute zur Definition von „kompakt“ verwendet wird. Ausführlich werden Begriffe der Relativtopologie besprochen, etwa „A ist abgeschlossen in M“. Von großer Wirkung ist Hausdorffs Definition des Zusammenhangs: „Eine von Null verschiedene Menge A heißt zusammenhängend, wenn sie sich nicht in zwei fremde, von Null verschiedene, in A abgeschlossene Teilmengen spalten läßt.“ Nach einer Behandlung des Begriffs „dicht“ und der in 2. 12 wiedergegebenen Konstruktion perfekter nirgendsdichter Mengen gibt Hausdorff dann zum Abschluß des Kapitels einige Anwendungen der Begriffe für die reellen Zahlen.

Hausdorff 1914 Grundzüge der Mengenlehre, Kapitel 8: Punktmengen in speziellen Räumen.

Die Themen in diesem fast 100 Seiten starken Kapitel sind:

- erstes und zweites Abzählbarkeitsaxiom (in heutiger Sprache: jeder Punkt hat eine abzählbare offene Umgebungsbasis bzw. der Raum selbst hat eine abzählbare offene Umgebungsbasis). Hausdorff bewegt sich schrittweise von allgemeinen topologischen Räumen startend zu spezielleren Räumen, bis hin zu den metrischen Räumen.
- Cantors Punktmengenanalyse. Hausdorff startet mit dem Satz, daß $\subseteq$ -auf- oder $\subseteq$ -absteigende Folgen von offenen oder abgeschlossenen Mengen in Räumen, die das zweite Abzählbarkeitsaxiom erfüllen, höchstens abzählbar sind, und diskutiert den Cantorsche Ableitungsprozeß. Es folgt ein wichtiger weiterer Schritt in der deskriptiven Mengenlehre: Hausdorff führt den Begriff des Residuums $\varphi(M)$ einer Menge M ein, $\varphi(M) = \psi(\psi(M))$, mit $\psi(M) = M' - M$. Er nennt $\varphi(M)$ eine „viel energischere“ Reduktion von M als die Bildung von $M \cap M'$. M ist *reduzibel*, falls die iterierte Anwendung der φ -Operation zur leeren Menge führt. Hausdorff charakterisiert die reduzierbaren Mengen als gewisse Vereinigungen von Differenzen abgeschlossener Mengen und weiter als diejenigen Mengen, die sowohl $\mathcal{F}_\sigma$ (d. h. eine abzählbare Vereinigung von abgeschlossenen Mengen) als auch $\mathcal{G}_\delta$ (d. h. ein abzählbarer Schnitt von offenen Mengen) sind. Vieles ist beim Schreiben des Buches brandneu, und einige Resultate finden sich nur in den Anhängen des Buches.

- Beispiele für topologische Räume, die den euklidischen Räumen „bisweilen recht verschieden“ sind, etwa unendlich dimensionale Räume und Funktionenräume.
- eingehende Untersuchung metrischer Räume. Insbesondere werden Borelmengen für metrische Räume diskutiert, die durch Abschluß der abgeschlossenen Mengen F und der offenen Mengen G unter den Operationen δ ($\mathcal{A}_\delta$ = Menge aller abzählbaren Schnitte von Mengen aus $\mathcal{A}$) und σ ($\mathcal{A}_\sigma$ = Menge aller abzählbaren Vereinigungen von Mengen aus $\mathcal{A}$) entstehen. Dies liefert eine Hierarchie der Länge ω_1 , deren erste Stufen $\mathcal{G}, \mathcal{F}, \mathcal{G}_\sigma, \mathcal{F}_\delta, \mathcal{G}_{\delta\sigma}, \dots$ sind. Die Borel-Hierarchie wird fortan zur unverzichtbaren Leiter zu den Aussichtspunkten der deskriptiven Mengenlehre.
- kompakte Mengen in metrischen Räumen, totale Beschränktheit, Bedingungen für Kompaktheit.
- vollständige Räume. Hausdorff beweist den Satz von Young (1903): In einem vollständigen metrischen Raum ist jede $\mathcal{G}_\delta$ -Menge, d. h. jeder abzählbare Schnitt von offenen Mengen, abzählbar oder von der Mächtigkeit von $\mathbb{R}$. Im Anhang beweist Hausdorff die Kontinuumshypothese sogar für alle $\mathcal{G}_{\delta\sigma}$ -Mengen in vollständigen metrischen Räumen, die eine abzählbare dichte Teilmenge besitzen (heute: polnische Räume). Dieses Resultat ist eine wichtige Vorstufe zum Beweis der Kontinuumshypothese für alle Borelmengen in polnischen Räumen (Hausdorff und unabhängig Alexandrov 1916).
- Euklidische Räume $\mathbb{R}^n$ mit der üblichen Metrik, Polygone und Wege in der Ebene $\mathbb{R}^2$.

Hausdorff 1914 Grundzüge der Mengenlehre, Kapitel 9:
Abbildungen oder Funktionen.

Nach verschiedenen Charakterisierungen von „fist stetig im Punkt a“ mit Hilfe des Umgebungsbegriffs beweist Hausdorff zentrale Sätze über stetige Funktionen, etwa daß der Zusammenhang unter dem Bild einer stetigen Funktion erhalten bleibt. Anschließend werden (Jordan-) Kurven und der Dimensionsbegriff besprochen, wobei auch die Hilbertsche Surjektion (vgl. Anhang 1. 9.) diskutiert wird. Schließlich werden noch Folgen von stetigen Funktionen betrachtet, nebst den Bedingungen, die einen stetigen Limes garantieren.

Hausdorff 1914 Grundzüge der Mengenlehre, Kapitel 10:
Inhalte von Punktmengen.

Hausdorffs Behandlung dieses weiten Feldes hier ebenfalls nur in Schlagworten: Inhaltsbegriff, Peano-Jordan-Inhalt, Lebesgue-Maß, Lebesgue-Integral. Von besonderer Bedeutung ist Hausdorffs paradoxe Zerlegung der Kugeloberfläche, die sich in einem Nachtrag zu Kapitel 10 ganz am Ende der Grundzüge findet: Hausdorff zerlegt die Oberfläche K der dreidimensionalen Einheitskugel U in vier paarweise disjunkte Mengen A, B, C und Q derart, daß Q abzählbar ist, und die drei Mengen A, B, C sowohl untereinander kongruent sind (d. h. in diesem Fall: durch Drehungen zur Deckung gebracht werden können), als auch kongruent zu $B \cup C$. Banach und Tarski haben in [Banach / Tarski 1924] ähnlich „paradoxe“ („wilde“) Zerlegungen für die Kugel selbst konstruiert. (Der Leser sei dahingehend beruhigt, daß sich die Zerlegungen von Hausdorff, Banach und Tarski weder mit der Laubsäge noch mit modernster Computertechnik simulieren lassen.)

Hausdorffs Zerlegung zeigt, daß es keinen kongruenzrespektierenden Inhalt auf K gibt (zur Definition eines Inhalts siehe 2. 9). Ein Inhalt I auf der Kugeloberfläche K *respektiert Kongruenz*, falls gilt:

$$(+)\quad I(A) = I(B) \text{ für kongruente } A, B \subseteq K,$$

Daß es nun einen Inhalt I mit (+) auf $\mathcal{P}(K)$ wegen der Hausdorff-Zerlegung nicht geben kann, sieht man so. Zunächst müßte $I(Q) = 0$ gelten: *Annahme*, $I(Q) = \delta > 0$. Sei dann $n \in \mathbb{N}$

derart, daß $n \cdot \delta > 1$ gilt. Da Q abzählbar ist, findet man leicht Drehungen $\pi_1, \dots, \pi_n$ der Kugel, die die Menge Q in paarweise disjunkte Mengen $Q_1, \dots, Q_n$ überführen. Dann ist aber $I(Q_1 \cup \dots \cup Q_n) = n \cdot I(Q) > 1$, *Widerspruch*. Also haben wir $I(Q) = 0$ und damit also $I(A \cup B \cup C) = 1$. Aber $I(A) = I(B) = I(C) = 1/3$, $I(A) = I(B \cup C) = I(B) + I(C) = 2/3$ nach Konstruktion von A, B, C , *Widerspruch*.

Hausdorff 1916 Die Mächtigkeit der Borelschen Mengen

Hausdorff zeigt, daß jede überabzählbare Borelmenge in $\mathbb{R}$ eine nichtleere perfekte Teilmenge besitzt. Damit ist jede Borelmenge abzählbar oder von der Mächtigkeit des Kontinuums. Unabhängig von Hausdorff hat diesen Satz im gleichen Jahr auch Alexandrov bewiesen.

Hausdorff 1927 Mengenlehre (zweite Auflage der „Grundzüge“)

„Mengenlehre“ ist eine sowohl gekürzte als auch erweiterte Auflage des Originals. Im Vorwort notiert Hausdorff hierzu, „daß der äußere Rahmen, in dem das Buch jetzt erscheint, eine starke Einschränkung des Umfangs gegenüber der ersten Auflage“ erforderte. Die Behandlung der Maß- und Integrationstheorie fällt weg, und statt dem „topologischen Standpunkt, durch den sich die erste Auflage anscheinend viele Freunde erworben“ hatte, werden nun fast ausschließlich nur noch metrische Räume betrachtet. Weiter wurde der Ordnungstheorie eine Diät auferlegt, dafür bringt das Buch aber neue Pralinen aus deskriptiver Sicht. Hausdorff diskutiert viel ausführlicher als in den „Grundzügen“ die Borel-Hierarchie, und kommt dann zu den Suslinschen Mengen, die heute analytische Mengen heißen, und bequem als die Bilder von Borelmengen unter stetigen Funktionen eingeführt werden können. Äquivalent ist die Definition über die sog. $\mathcal{A}$ -Operation, die (wahrscheinlich) auf Suslin und Alexandrov zurückgeht (zur Definition von $\mathcal{A}$ siehe 2.12). Die $\mathcal{A}$ -Operation, angewendet auf abgeschlossene Mengen, erzeugt genau die analytischen Mengen. Bemerkenswert ist, daß man in einem Schritt nur aus den abgeschlossenen Mengen sogar mehr als die Borelmengen erhält. Es gibt analytische Mengen, die nicht Borelsch sind, und die Borelmengen sind genau die analytischen Mengen A , deren Komplement $\mathbb{R} - A$ ebenfalls analytisch ist (Satz von Suslin 1917). Weiter kann man die Kontinuumshypothese (und stärker die Scheffer-Eigenschaft) für die analytischen Mengen zeigen (Lusin, Suslin 1917).

All diese Dinge über „einfache, definierbare Teilmengen von $\mathbb{R}$ “, und mehr, finden sich im Buch von Hausdorff, und einige Paragraphen zusammengekommen bilden für sich ein kleines Lehrbuch der elementaren deskriptiven Mengenlehre.

Zermelo 1901 Über die Addition transfiniten Kardinalzahlen

In seiner ersten mengentheoretischen Arbeit beweist Zermelo: Sind α und b_n , $n \in \mathbb{N}$, Kardinalzahlen, und gilt $\alpha + b_n = \alpha$ für alle $n \in \mathbb{N}$, so gilt $\alpha + \sum_{n \in \mathbb{N}} b_n = \alpha$. Insbesondere folgt aus $\alpha + b = \alpha$ für Kardinalzahlen α und b , daß $\alpha = \alpha + \omega \cdot b$ gilt (Fall $b = b_n$ für alle n). 1901 standen weder der Wohlordnungssatz noch der Multiplikationssatz zur Verfügung. Kardinalzahlen werden hier im Sinne von Cantors Definition einer Kardinalzahl durch Abstraktion angesehen. Der wesentliche Trick der Argumentation ist die Verwendung der trivialen Gleichung $2 \cdot \omega = \omega$, und besteht in der folgenden Rechnung, die den „insbesondere“-Teil liefert: Sei $\alpha = \alpha + b$. Dann gilt $\alpha = \alpha + b + b = \alpha + b + b + b$, usw. Das „usw.“ liefert ein $c \geq 0$ mit $\alpha = \omega \cdot b + c$ (!). Dann gilt aber $\alpha = 2 \cdot \omega \cdot b + c = \omega \cdot b + c + \omega \cdot b = \alpha + \omega \cdot b$.

Zermelo hält fest, daß sich sein Beweis als eine Verallgemeinerung von Bernsteins Beweis des Satzes von Cantor-Bernstein auffassen läßt: Die ω -vielen b 's in $\alpha = b + b + \dots + c$ entsprechen den „Orbitmengen“ $C_0, C_1, \dots$ im Beweis von Cantor-Bernstein (mit $|C_0| = b$).

Zermelo 1904 Beweis, daß jede Menge wohlgeordnet werden kann

Untertitel: „Aus einem an Herrn Hilbert gerichteten Briefe.“ (Für Auszüge dieser Arbeit siehe 2.5.)

In den drei Seiten dieser Arbeit zeigt Zermelo seinen Wohlordnungssatz: „Jede Menge läßt sich wohlordnen.“ Die Arbeit setzt eine kontroverse Diskussion über das im Beweis verwendete (und deutlich als unbeweisbares Prinzip isolierte) Auswahlaxiom in Gang.

Die insbesondere von Hilbert in seiner Problemsammlung 1900 geäußerte Hoffnung, mit Hilfe einer Wohlordnung der reellen Zahlen das Kontinuumsproblem lösen zu können, erfüllt sich nicht: Zermelos Beweis gibt keine Auskunft über die Länge einer solchen Wohlordnung, anhand derer man die Mächtigkeit der reellen Zahlen ablesen könnte.

Als Konsequenzen des Wohlordnungssatzes hebt Zermelo den Vergleichbarkeitssatz und den Multiplikationssatz hervor.

Zermelo 1908 Neuer Beweis für die Möglichkeit einer Wohlordnung

Der erste Teil der Arbeit (§1) bringt einen gegenüber dem Beweis von 1904 anders organisierten Beweis des Wohlordnungssatzes; Zermelo vermeidet hier jeden Bezug auf die Ordnungstheorie, insbesondere wird der Fundamentalsatz für Wohlordnungen nicht mehr benutzt. Wieder wird die Verwendung des Auswahlaxioms klar dargestellt.

Der Hauptteil der Arbeit (§2) setzt sich dann (polemisch, aber inhaltlich korrekt) mit den vielen Kritiken an seinem Wohlordnungssatz auseinander. Zusammenfassend betont Zermelo, daß seit dem Jahre 1904 „trotz eingehender Prüfung“ keine mathematischen Fehler in seinem Argument gefunden wurden.

Zermelo 1908 Untersuchungen über die Grundlagen der Mengenlehre. I

Diese Arbeit enthält die mengentheoretische Axiomatik von Zermelo. Ziel ist, „daß man die Prinzipien [der Mengenbildung] einmal eng genug einschränkt, um alle Widersprüche auszuschließen, gleichzeitig aber auch weit genug ausdehnt, um alles Wertvolle beizubehalten.“ Zermelo spricht von einem „Bereich“ von Objekten [mengentheoretisches Universum]. Die Objekte, die er auch „Dinge“ nennt, zerfallen in Mengen und Urelemente. Zwischen den Dingen gibt es als Grundrelation die Elementbeziehung. In dem Bereich aller Objekte gelten dann die sieben Zermeloschen Axiome: Axiom der Bestimmtheit [Extensionalitätsaxiom], der Elementarmengen, der Aussonderung, der Potenzmenge, der Vereinigung, der Auswahl, des Unendlichen. Die Russell-Zermelo-Antinomie fällt weg, denn aus den Axiomen folgt: „Jede Menge M besitzt mindestens eine Untermenge [Teilmenge] M_0 , die nicht Element von M ist.“ Hieraus folgt, „daß nicht alle Dinge x des Bereiches $\mathfrak{B}$ Elemente ein und derselben Menge sein können; d. h. der Bereich $\mathfrak{B}$ ist selbst keine Menge.“

Der zweite Teil der Arbeit behandelt den Mächtigkeitsbegriff. Die Gleichmächtigkeit zweier Mengen M, N definiert Zermelo – mangels des Begriffs eines geordneten Paares – zunächst nur zwischen disjunkten Mengen. Sind M, N disjunkt, so ist eine Abbildung [Bijektion] von M nach N definiert als eine Menge Φ bestehend aus ungeordneten Paaren $\{m, n\}$ mit $m \in M, n \in N$, derart, daß jedes Element von $M \cup N$ in genau einem $\{m, n\} \in \Phi$ vorkommt. Disjunkte M und N heißen gleichmächtig, falls eine Bijektion von M nach N existiert. Später wird dann auch die Gleichmächtigkeit beliebiger Mengen M, N definiert: M und N heißen gleichmächtig, falls eine zu M und N disjunkte Menge R existiert, die sowohl gleichmächtig zu M als auch zu N ist.

Es folgen die Hauptsätze der Theorie der Mächtigkeiten; insbesondere zeigt Zermelo den Äquivalenzsatz [Satz von Cantor-Bernstein] ohne Verwendung der natürlichen Zahlen – mit Referenz an die Dedekindsche Kettentheorie.

Weiter wird der Satz von Cantor und der Satz von König-Zermelo aus der Kardinalzahlarithmetik bewiesen.

Zermelo 1930 Über Grenzzahlen und Mengenbereiche: Neue Untersuchungen über die Grundlagen der Mengenlehre

Zermelo stellt seine abschließende Axiomatik vor und untersucht Modelle für dieses Axiomensystem. Er nimmt eine nichtformalistische Haltung ein (und entwickelt sein System weiterhin fern der Prädikatenlogik erster Stufe).

Seine Axiomatik umfaßt nun das Fraenkelsche Ersetzungsaxiom und das Fundierungsaxiom. Andererseits betrachtet er das Auswahlaxiom hier nicht als Axiom, sondern als „allgemeines logisches Prinzip“. Ebenso fehlt das Unendlichkeitsaxiom „als nicht zur allgemeinen Mengenlehre gehörig“. Modelle, in denen das Unendlichkeitsaxiom falsch ist, betrachtet er aber „nicht als wahre ‚Modelle‘ der Cantorsche Mengenlehre“. Er bemerkt: „Aus ihnen heraus führt erst mein früheres Axiom des Unendlichen.“

Das Fundierungsaxiom hat Zermelo unabhängig von Fraenkel und von Neumann formuliert. Die Arbeit enthält bereits die Idee zum Beweis der Widerspruchsfreiheit dieses Axioms relativ zu den übrigen Axiomen durch Rückzug auf den fundierten Teil des Universums.

Zermelo untersucht „Bereiche“, bestehend aus Mengen und Urelementen, „in denen diese Axiome gelten, sogenannte „Normalbereiche“. Zermelo zeigt, daß Normalbereiche durch zwei Zahlen bis auf Isomorphie eindeutig bestimmt sind: Die Mächtigkeit der Urelemente und ihre Höhe, d. h. den Typus ihrer Ordinalzahlen.

Wesentliches Hilfsmittel der Untersuchung ist die „Schichtung“ eines Normalbereichs. Zermelo rückt hier zum ersten Mal die kumulative Hierarchie in den Mittelpunkt, die zuvor schon von Mirimanov (1917) und von Neumann (1925) angedeutet wurde.

In heutiger Sprache kann man sein Vorgehen und seine Resultate so beschreiben: Zermelo bildet $V_0 =$ „alle Urelemente“, $V_{\alpha+1} = \mathcal{P}(V_\alpha) \cup V_0$ für alle Ordinalzahlen α , $V_\lambda = \bigcup_{\alpha < \lambda} V_\alpha$ für alle Limesordinalzahlen λ . Er zeigt dann, daß V_α genau dann ein Normalbereich ist, wenn α eine stark unerreichbare Kardinalzahl („Grenzzahl“) ist, wobei „Normalbereich“, d. h. Modell für die Axiome, zweitstufig interpretiert wird (und werden muß für die Beweisrichtung „ V_α ist Normalbereich“ folgt „ α stark unerreichbar“).

Zermelo betont, daß er die Existenz von Normalbereichen nicht aus den Axiomen ableiten kann, entwickelt aber ein Bild der Mengenwelt, in der ein Normalbereich auf den anderen folgt und diesen fortsetzt, sodaß die echten Klassen des einen Bereiches zu Mengen in den darauffolgenden Bereichen werden. Insbesondere kann dann jeder Normalbereich auch als Menge aufgefaßt werden. Die Folge der derart aufeinander aufbauenden Normalbereiche betrachtet Zermelo analog zu den Ordinalzahlen als unbegrenzt. Konsequenter formuliert er „die Existenz einer unbegrenzten Folge von Grenzzahlen [stark unerreichbaren Kardinalzahlen]“ als „neues Axiom für die ‚Meta-Mengenlehre“.

4. Zeittafel zur frühen Mengenlehre

Die folgende Tafel dokumentiert wichtige Stationen der Entwicklung der Mengenlehre. Abgedeckt ist etwa der Zeitraum von 1873 – 1930, wobei darüber hinaus die mengentheoretischen Resultate von Gödel aus den 30er Jahren und Cohen von 1963 angesprochen werden. Die Tafel endet mit einem sehr knappen Eintrag zu modernen Entwicklungen.

Daten zur allgemeinen Logik sind nicht eingearbeitet. Zur Auswahl ist zu darüber hinaus zu sagen, daß der Autor unter „Mengenlehre“ die (intuitiv oder formal präsentierte) mathematische Theorie versteht, in deren Zentrum die Begriffe der Mächtigkeit und der Wohlordnung stehen. Daher sind keine Daten aufgenommen, die Beiträge zur bloßen Bildung des Konzepts „Menge“ – etwa solche von Bolzano, Dedekind, Riemann – und seiner wachsenden Bedeutung für die Mathematik des 19. Jahrhundert widerspiegeln würden.

Die Jahreszahlen beziehen sich auf den Zeitpunkt der Veröffentlichung von Resultaten, wenn es aus dem kommentierenden Text nicht anders hervorgeht.

7. Dezember 1873 : Überabzählbarkeit von $\mathbb{R}$

Geburtstag der Mengenlehre. Cantor beweist die Überabzählbarkeit der reellen Zahlen, wobei die Vollständigkeit von $\mathbb{R}$ benutzt wird – das wohl in den 80er Jahren gefundene Diagonalargument stellt er erst 1891 der Öffentlichkeit vor. Er veröffentlicht sein Resultat 1874 zusammen mit einem auch von Dedekind gefundenen Beweis der Abzählbarkeit der algebraischen Zahlen. Es ergibt sich so ein neuer Beweis für die Existenz transzendenter Zahlen. (Manche Historiker würden das „auch“ im vorletzten Satz wahrscheinlich streichen. Cantor zumindest hat Dedekind in der Arbeit unerwähnt gelassen, und später aufrecht erhalten, den Beweis selbst gefunden zu haben. Der Autor dieses Buches bleibt bei „auch“.)

1878 : Mächtigkeitsbegriff, $|\mathbb{R}^n| = |\mathbb{R}|$, Kontinuumshypothese

Definition von „zwei Mengen haben die gleiche Mächtigkeit“ und „eine Menge hat kleinere Mächtigkeit als eine andere“ durch Cantor. Die Vergleichbarkeit von zwei Mengen bzgl. ihrer Mächtigkeit nimmt Cantor hier noch als gegeben an; das Problem wird erst 1904 durch den Wohlordnungssatz von Zermelo vollständig gelöst.

Cantor zeigt, daß $\mathbb{R}^n$ und $\mathbb{R}$ gleichmächtig sind für alle $n \geq 1$.

Erste Formulierung der Kontinuumshypothese.

Diagonalaufzählung von $\mathbb{N}^2$ durch ein Polynom zweiten Grades.

1880 : iterierte Ableitungen, arithmetische Notationen für die Stufen transfiniter Prozesse

Cantor macht die ersten Schritte im Reich der transfiniten Zahlen. Durch die iterierte Ableitung von Punktmengen entstehen erste arithmetische Notationen für die Stufen transfiniter Prozesse (in heutiger Sprache bis etwa hin zu ϵ_0). Cantor spricht von einer

„dialektische[n] Begriffserzeugung, welche immer weiter führt und dabei frei von jeglicher Willkür in sich notwendig und konsequent bleibt.“ In einer Fußnote gibt er an, mit dieser Begriffserzeugung der transfiniten Ordnungen schon seit zehn Jahren vertraut zu sein.

1882 : Klassisch-Platonischer Mengenbegriff, Separabilität von $\mathbb{R}^n$

Cantor nennt eine Menge wohldefiniert, „wenn auf Grund ihrer Definition und in Folge des logischen Prinzips vom ausgeschlossenen Dritten“ es „intern bestimmt“ ist, ob irgendein Objekt (einer gegebenen „Begriffssphäre“) der Menge angehört oder nicht, und ob zwei in verschiedener Weise definierte Elemente einer Menge identisch sind oder verschieden. Es ist für die Wohldefiniertheit und damit für die Existenz der Menge nicht wesentlich, ob das Elementsein von Objekten und das Gleichsein von Elementen tatsächlich auch „extern bestimmt“ ist, d. h. durch einen Beweis erkannt werden kann.

Beweis (in heutiger Formulierung), daß $\mathbb{R}^n$ für jedes $n \geq 1$ separabel ist, d. h. es existiert eine abzählbare dichte Teilmenge von $\mathbb{R}^n$.

1883 : Wohlordnungsbegriff, Erzeugungsprinzipien, perfekte Mengen, Kardinalzahlen und Ordinalzahlen als Universalien

Cantor beschreibt zwei „Erzeugungsprinzipien“ (Nachfolgerbildung und Limesbildung), die transfinite Zahlen generieren. ω -Notation für die erste transfinite Zahl.

Cantor definiert den Begriff einer wohlgeordneten Menge, und stellt damit dem Begriff der Mächtigkeit die wichtigste Begriffsbildung der Mengenlehre an die Seite.

Definition einer perfekten Menge von reellen Zahlen. Cantor gibt auch ein Beispiel für eine nirgends dichte nichtleere perfekte Menge – die heutige Cantormenge, deren fernöstlich anmutende Visualisierung der Leser auf dem Buchdeckel findet.

In einem Vortrag im September in Freiburg stellt Cantor der Öffentlichkeit seine Vorstellung von Kardinalzahlen und Ordinalzahlen als Allgemeinbegriff oder „Universalien“ vor, die er später in seinen Schriften wiederholt und erläutert (vgl. die Zitate am Ende von 1.4 und am Anfang von 2.6, und siehe [Cantor 1887].)

1884 : Lösung des Kontinuumsproblems für abgeschlossene Mengen

Definition einer abgeschlossenen Menge von reellen Zahlen durch Cantor. Satz von Cantor-Bendixson: Es existiert eine eindeutige Zerlegung einer abgeschlossenen Punktmenge in ihren perfekten Kern und einen abzählbaren Rest. Cantor zeigt weiter, daß nichtleere perfekte Mengen immer die Mächtigkeit der reellen Zahlen besitzen. Zusammengekommen ergibt sich, daß abgeschlossene Mengen immer abzählbar oder von der Mächtigkeit des Kontinuums sind, und damit ist die Kontinuumshypothese für die abgeschlossenen Mengen bewiesen.

1887 : Philosophische Arbeit von Cantor

In seinen „Mitteilungen zur Lehre vom Transfiniten. I.“ verteidigt Cantor seinen aktual unendlichen Mengenbegriff gegen philosophische Einwände, und er erläutert seine Vorstellung von Kardinal- und Ordinalzahlen als durch Abstraktion gewonnene Gebilde aus Einsen (vgl. Eintrag zu 1883).

1888 : Endlichkeitsdefinition von Richard Dedekind

In seinem Buch „Was sind und was sollen die Zahlen“ definiert Dedekind: Eine Menge ist *unendlich*, wenn sie zu einer echten Teilmenge von sich selbst gleichmächtig ist. Dieses Charakteristikum unendlicher Mengen findet sich bereits in den einleitenden Worten von [Cantor 1878].

1891 : Satz von Cantor: Das Diagonalargument

In einem Vortrag auf der ersten Tagung der Deutschen Mathematiker-Vereinigung in Halle zeigt Cantor durch Diagonalisierung, daß (in heutiger Formulierung) die Potenzmenge einer Menge immer größere Mächtigkeit hat als die Menge selbst (Satz von Cantor). Aus dem Beweis ergibt sich das Diagonalverfahren für eine Liste reeller Zahlen in Dezimaldarstellung, mit dem die Überabzählbarkeit von $\mathbb{R}$ heute zumeist bewiesen wird.

1895/1897 : Cantors abschließende systematische Darstellung der Mengenlehre

Zwei längere, lehrbuchartige Artikel von Cantor. (Neue Hauptresultate und Konstruktionen: Charakterisierung der Ordnungstypen von $\mathbb{Q}$ und $\mathbb{R}$, Vergleichbarkeitssatz für Wohlordnungen, Definition der Ordinalzahlpotenz durch transfinite Rekursion, Cantorsche Normalform, ε -Zahlen.) Cantor hat seit längerer Zeit Einsichten in „zu große“ Zusammenfassungen, veröffentlicht aber hierzu nichts.

1897 : Bernstein beweist den Satz von Cantor-Bernstein-Dedekind

In einem Seminar in Halle beweist der 19jährige Felix Bernstein den Satz. Cantor teilt den Beweis Borel auf dem ersten internationalen Mathematiker-Kongreß in Zürich mit, der ihn in [Borel 1898] veröffentlicht (als Resultat von Bernstein). Dedekind hatte 1887 einen Beweis gefunden, ihn aber nicht in seinem Buch „Was sind und was sollen die Zahlen?“ von 1888 veröffentlicht, und ihn auch Cantor nicht mitgeteilt. Zermelo hat 1908 den Beweis von Dedekind wiederentdeckt, dieser selbst findet sich im Nachlaß von Dedekind (siehe [Dedekind 1930 – 1932]).

1900 : Hilberts Problemliste, Bericht von Arthur Schoenflies

Hilbert setzt Cantors Kontinuumsproblem an die erste Stelle seiner Liste von offenen Fragen, die er auf dem zweiten internationalen Mathematiker-Kongreß in Paris vorstellt. Arthur Schoenflies verfaßt im Auftrag der Deutschen Mathematiker-Vereinigung einen 250-seitigen Bericht über „Die Entwicklung der Lehre von den Punktmannigfaltigkeiten“. Dieser Bericht wird 1908 ergänzt und erscheint 1913 in überarbeiteter Form.

1904 : Zermelos Wohlordnungssatz

Zermelo beweist den Wohlordnungssatz. Sein Beweis stößt auf zum Teil heftige Ablehnung aufgrund der Verwendung eines abstrakten Auswahlprinzips.

1905 : Ungleichung von Julius König

König zeigt den heute nach ihm (und Zermelo) benannten Satz der Kardinalzahlarithmetik.

1906 : Hessensbergs Buch „Grundbegriffe der Mengenlehre“, Beweis des Multiplikationssatzes

Hausdorff nennt in seinem eigenen Lehrbuch von 1914 Hessensbergs Buch „die erste allgemein gehaltene Darstellung“ der Theorie der wohlgeordneten Mengen, und das Buch ist bis 1914 sicher die beste Einführungsliteratur in die Mengenlehre neben den Cantorschen Originalarbeiten. Darüber hinaus enthält es neue Ergebnisse und Begriffe. Hessenberg untersucht parallel zu und unabhängig von Hausdorff den Konfinalitätsbegriff, und weiter betrachtet er eingehend Fixpunkte von Normalfunktionen (in heutiger Sprache), in Fortsetzung der Theorie der ϵ -Zahlen von Georg Cantor.

In seinem Buch beweist Hessenberg auch zum ersten Mal vollständig den Multiplikationssatz: Es gilt $|M \times M| = |M|$ für alle unendlichen Mengen M . 1908 legt Jourdain einen zweiten Beweis vor. Hausdorff hatte schon 1904 die Regularität von Nachfolgerkardinalzahlen behauptet, aus der der Satz leicht folgt. Bernstein berichtet in seiner Dissertation 1901, daß bereits Cantor den Satz für wohlordenbare Mengen bewiesen hatte.

1906 – 1908 : Untersuchung von linearen Ordnungen durch Hausdorff

Felix Hausdorff arbeitet völlig unbeeindruckt von den Paradoxien und der Grundlagediskussion an einer allgemeinen Theorie der Ordnungstypen von linear geordneten Mengen, in Fortsetzung der Cantorschen Untersuchung von Ordinalzahlen und ihrer Arithmetik.

1908 : Zermelo-Axiomatik, Existenz von Bernstein-Mengen, Frage nach der Existenz unendlich großer Zahlen

Zermelo gibt einen zweiten Beweis für den Wohlordnungssatz und stellt ein nichtformales Axiomensystem auf, welches die Prinzipien isoliert, die für eine Entwicklung der Grundbegriffe der Mengenlehre und einen Beweis des Wohlordnungssatzes gebraucht werden. Das System enthält bis auf Fundierung und Ersetzung alle Axiome der Zermelo-Fraenkel-Axiomatik. Es dauert bis nach dem zweiten Weltkrieg, bis das System von Zermelo vervollständigt und in der Sprache der Prädikatenlogik formuliert wird, und in dieser Form dann als geeignete und genügend gesicherte Grundlagentheorie für die Mathematik allgemeine Akzeptanz findet.

Bernstein zeigt, daß es eine Menge reeller Zahlen B gibt derart, daß weder die Menge B noch ihr Komplement $\mathbb{R} - B$ eine nichtleere perfekte Teilmenge besitzen.

Hausdorff diskutiert die Frage, ob es reguläre Kardinalzahlen $\aleph_\alpha$ gibt derart, daß α eine Limeszahl ist. Heute heißen solche Zahlen schwach unendlich große Kardinalzahlen. (Eine Kardinalzahl κ ist dabei regulär, wenn jede Teilmenge von $W(\kappa)$, deren Mächtigkeit kleiner als κ ist, in κ beschränkt ist.)

1909 : Hessenberg über Kettentheorie

Gerhard Hessenberg verwendet die Zermelo-Axiomatik in einer Arbeit über Kettentheorie und Wohlordnung.

1911 : Paul Mahlo untersucht große Kardinalzahlaxiome

Aufbauend auf der Arbeit von Hausdorff von 1908 untersucht Mahlo zwischen 1911 und 1913 unerreichbare Kardinalzahlen, unerreichbare Kardinalzahlen, die ein Limes von unerreichbaren Kardinalzahlen sind, usw. Er findet auch den nach diesem „usw.“ liegenden, heute nach ihm benannten Begriff einer Mahlo-Kardinalzahl.

1914 : Lehrbuch „Grundzüge der Mengenlehre“ von Felix Hausdorff

In den Grundzügen gibt Hausdorff eine systematische und nichtaxiomatische Einführung in den aktuellen Stand der Cantorschen Mengenlehre und seine eigene Theorie der linearen Ordnungstypen. Gleiches Gewicht kommt dann der Untersuchung von „Punktmengen in allgemeinen Räumen“ zu. Hausdorff definiert topologische und metrische Räume und diskutiert Grundlagen der Maß- und Integrationstheorie. Der Text ist ein Jahrhundertwerk der Lehrbuchliteratur mit großer, wenn auch gemächlicher Wirkung. Letztendlich ist es der topologische Teil des Buches, der die größte Beachtung erfährt.

1915 : Äquivalenz von Vergleichbarkeitssatz und Wohlordnungssatz

Friedrich Hartogs zeigt, daß der Vergleichbarkeitssatz und der Wohlordnungssatz über den Axiomen von Zermelo ohne Auswahlaxiom äquivalent sind. Die Arbeit ist eine der wenigen, die die Axiomatik von Zermelo aufgreifen.

1916 : Lösung des Kontinuumsproblems für die Borelmengen

Hausdorff und Alexandrov beweisen unabhängig voneinander: Jede Borelmenge hat die Scheeffer-Eigenschaft (ist abzählbar oder besitzt eine nichtleere perfekte Teilmenge). Damit ist jede Borelmenge von reellen Zahlen abzählbar oder gleichmächtig mit $\mathbb{R}$. Vorstufen dieses Resultats sind: Cantor 1884 für abgeschlossene Mengen; Young 1903 für $\mathcal{G}_\delta$ -Mengen, und hierauf aufbauend Hausdorff 1914 für $\mathcal{G}_{\delta\sigma\delta}$ -Mengen. 1917 zeigt die russische Schule um Alexandrov, Suslin und Lusin dann die Scheeffer-Eigenschaft sogar für die analytischen Mengen, d. h. die Bilder von Borelmengen unter stetigen Funktionen. (Bereits für die Menge der Komplemente von analytischen Mengen ist die Scheeffer-Eigenschaft dann unabhängig von ZFC. Insbesondere gibt es in Gödels Modell L eine analytische Menge A , sodaß $\mathbb{R} - A$ überabzählbar ist, aber keine perfekte Teilmenge besitzt. Genügend große Kardinalzahlen garantieren dagegen die Scheeffer-Eigenschaft für diese und viele weitere „einfache“ Mengen von reellen Zahlen.)

1919 : Lehrbuch von Abraham Fraenkel

Fraenkels „Einleitung in die Mengenlehre“ erscheint, basierend auf „Unterhaltungen, in denen ich Kriegskameraden (Nichtmathematikern) gelegentlich öde Stunden durch Einführung in Gedankengänge der Mengenlehre verkürzen konnte.“ 1923 und 1928 erscheinen erweiterte Auflagen, und Fraenkels Einleitung begleitet in dieser Zeit das Buch von Hausdorff als ein Standardwerk zur Mengenlehre. Die Mathematik ist hier weitgehend elementar, jedoch behandelt das Buch philosophische Aspekte und die Grundlagen Diskussion.

1922 : Skolems Kritik, Ersetzungsaxiom, Fundierung

In einem Vortrag auf einem Kongreß in Helsingfors kritisiert Thoralf Skolem Zermelos Begriff einer „definiten Aussage“ in dessen Aussonderungsaxiom, und schlägt (in heutiger Sprache) eine Formulierung des Eigenschaftsbegriffs in der Prädikatenlogik erster Stufe vor. Zermelo hat eine solche Formalisierung seiner Axiomatik abgelehnt [vgl. Zermelo 1929]. Es dauert dann auch noch etliche Jahre, bis die Formalisierung vollständig durchgeführt und akzeptiert ist.

Fraenkel diskutiert das Ersetzungsaxiom; ihm war 1921 aufgefallen, daß in der Zermelo-Mengenlehre die Existenz der Kardinalzahl $\aleph_\omega$ nicht beweisbar ist. Unabhängig davon war ein Ersetzungsaxiom auch von Skolem 1922 und zuvor von Mirimanov 1917 betrachtet worden. Fraenkel schlägt außerdem ein Axiom der Beschränkung vor, dessen exakte Formulierung jedoch Schwierigkeiten macht. In eine exakte Form brachten es von Neumann 1925 und unabhängig hiervon Zermelo 1930. Zermelo benannte das Axiom als Fundierungsaxiom.

1923 : Ordinalzahldefinition von John von Neumann

John von Neumann definiert kanonische Repräsentanten für alle Klassen bestehend aus Wohlordnungen gleicher Länge. Die Konstruktion ist formal in der um das Ersetzungsaxiom angereicherten Zermelo-Axiomatik durchführbar, und die definierten Repräsentanten sind dann die (Neumann-Zermelo-) Ordinalzahlen, wie sie heute verwendet werden. (Das Ersetzungsaxiom wird gebraucht, um einen Repräsentanten für *alle* Klassen zu erhalten; in Modellen der Zermelo-Mengenlehre gibt es im allgemeinen nur wenige Neumann-Zermelo-Ordinalzahlen.)

1925 – 1928 : Axiomatisierung von John von Neumann

Von Neumann entwickelt eine eigene Axiomatik, die den Funktionsbegriff statt den Mengenbegriff als Grundbegriff verwendet. In den 30er Jahren wurde diese Axiomatik von Bernays und dann noch einmal von Gödel 1940 vereinfacht, und dann mit Hilfe des Mengenbegriffs formuliert. Das resultierende System ist heute als NBG = Neumann-Bernays-Gödel-Mengenlehre bekannt. In ihr gibt es echte Klassen als Objekte, die Theorie NBG beweist aber dieselben Sätze über Mengen wie ZFC. Lange lief die Theorie NBG gleichbeachtet neben ZFC, ab den 50er Jahren wurde ZFC zur Standardtheorie.

1929 / 1930 : kumulative Hierarchie, Zermelo-Fraenkel-Axiomatik

Von Neumann untersucht 1929 die kumulative Hierarchie aller Mengen in seiner Axiomatik, Zermelo tut dies 1930 in einer Arbeit, in der er seine Axiome um das Ersetzungsaxiom und das Fundierungsaxiom erweitert. Die so entstehende Axiomatik wird seither als Zermelo-Fraenkel-Axiomatik bezeichnet, und wird heute mit ZFC abgekürzt (wobei dann zumeist die formalisierte Version gemeint ist).

1929 : Ultrafiltersatz von Tarski-Ulam

Ulam zeigt, daß ein nichttrivialer Ultrafilter auf $\mathbb{N}$ existiert. Tarski verallgemeinert dieses Ergebnis: Jeder Filter auf einer Menge M läßt sich zu einem Ultrafilter erweitern. Ulams Beweis verwendet eine Wohlordnung von $\mathbb{R}$, Tarskis Beweis das Auswahlaxiom.

1930 : Ulam untersucht vollständige Maße Buch über deskriptive Mengenlehre von Lusin

Ulam untersucht Maße auf Kardinalzahlen mit starken Vollständigkeitseigenschaften, was zum Begriff der meßbaren Kardinalzahl führt. William Hanf und Alfred Tarski zeigen dann Anfang der 60er Jahre, daß die Existenz meßbarer Kardinalzahlen viel stärker ist als die Existenz von unerreichbaren Kardinalzahlen oder auch von Mahlo-Kardinalzahlen.

Von Nikolai Lusin erscheint ein Buch über deskriptive Mengenlehre. Alexandrov, Lusin, Suslin und Sierpiński hatten ab 1916 die deskriptive Mengenlehre weiterentwickelt (vgl. hierzu auch den Eintrag in 3.3 zum Buch von Hausdorff 1927). Insbesondere wurden projektive Mengen auf ihre Regularitätseigenschaften (Lebesgue-Meßbarkeit, Baire-Eigenschaft, Scheeffe-Eigenschaft, Uniformisierbarkeit) untersucht. Die projektiven Teilmengen von $\mathbb{R}$ sind dabei diejenigen Punktmengen, die man aus einer abgeschlossenen Menge eines Raumes $\mathbb{R}^n$ durch iterierte Anwendung der Operationen „Projektion auf kleinere Dimensionen“ und „Komplementbildung“ erhalten kann. (Alle Borelmengen und alle analytischen Mengen sind projektiv; die analytischen Teilmengen von $\mathbb{R}$ sind dabei genau die Projektionen von $\mathcal{G}_\delta$ -Mengen im $\mathbb{R}^2$, und bilden zusammen mit ihren Komplementen das Anfangsstück einer natürlichen Hierarchie von projektiven Mengen.) Heute kennt man ein „Super-Regularitätsaxiom“ für die projektiven Mengen – das Axiom der projektiven Determiniertheit –, und weiß, daß es aus der Existenz von großen Kardinalzahlen folgt. Dinge weit draußen im Universum ermöglichen Beweise über definierbare Teilmengen von $\mathbb{R}$.

1935 : Relative Widerspruchsfreiheit des Auswahlaxioms

Kurt Gödel beweist die relative Konsistenz des Auswahlaxioms über der Theorie ZF (= ZFC ohne Auswahlaxiom). Das heißt genau: Ist ZF widerspruchsfrei, so ist auch ZFC widerspruchsfrei. Anders formuliert: Das Auswahlaxiom ist nicht verantwortlich, wenn in ZFC Widersprüche auftreten sollten. Gödel arbeitet in der Theorie ZF und definiert dort das sog. konstruktible Universum L , eine Teilklasse des Universums V . Die Klasse L enthält alle Ordinalzahlen, darüber hinaus aber nur die Mengen, die von den Axiomen explizit gefordert werden. Es zeigt sich, daß L ein Modell von ZF ist, und daß zudem L eine definierbare Wohlordnung besitzt. Damit ist L ein Modell von ZFC, und das Auswahlaxiom gilt dort in einer sehr starken globalen Form. Gödel veröffentlicht das Ergebnis 1938 zusammen mit dem analogen Ergebnis über die Kontinuumshypothese von 1937.

1937 : Relative Widerspruchsfreiheit der allgemeinen Kontinuumshypothese

Gödel zeigt, daß in seinem Modell L die verallgemeinerte Kontinuumshypothese richtig ist. In L läßt sich die Konstruktion von Teilmengen von ω genau unter die Lupe nehmen, und es zeigt sich, daß es genau ω_1 -viele Schritte im Aufbau von L gibt, bei denen eine neue Teilmenge von ω zu L hinzukommt. Also gilt $|\mathcal{P}(\omega)| = \omega_1$ in L , d. h. die Kontinuumshypothese gilt im konstruktiblen Universum. Die gleiche Beweisidee liefert die Gültigkeit der verallgemeinerten Kontinuumshypothese. (CH) und stärker (GCH) können also zu ZFC hinzugenommen werden, ohne daß die Widerspruchsfreiheit von ZFC dabei zerstört wird.

Die beiden Konstruktionen von Gödel zeigen: Das natürlichste Modell, das in ZF konstruiert werden kann, ist ein Modell des Auswahlaxioms und der verallgemeinerten Kontinuumshypothese. Ein Korollar zum Beweis ist: ZFC kann nicht beweisen, daß V und L verschieden sind (immer vorausgesetzt, daß ZFC konsistent ist.)

1938 – 1940 : Gödel veröffentlicht seine Resultate über L

1938 erscheint eine Zusammenfassung (für NBG), 1939 eine Beweisskizze (für ZF) und 1940 ein vollständiger Beweis der Resultate (für NBG).

Gödel hat weiter früh gesehen, daß L für die deskriptive Mengenlehre unangenehme Eigenschaften hat: Es gibt dort einfache projektive Mengen, denen die Regularitätseigenschaften nicht zukommen [vgl. Gödel 1938].

1963 : Die Erzwingungsmethode von Paul Cohen

Paul Cohen erfindet seine Erzwingungsmethode (forcing), ein sehr flexibles Werkzeug zur Konstruktion von Modellen. Während Gödels Konstruktion sich nach innen richtet und eine Teilklasse von V definiert, ist die Methode von Cohen extrovertiert und erweitert gegebene Modelle durch kontrolliertes Hinzufügen neuer Elemente. Cohen kann mit dieser Methode Modelle konstruieren, in denen die Kontinuumshypothese verletzt ist, und ein zusätzlicher Trick ermöglicht die Konstruktion eines Modells von ZF, in dem das Auswahlaxiom nicht gilt. Zusammen mit den Resultaten von Gödel ergibt sich so die Unabhängigkeit des Auswahlaxioms über der Theorie ZF und die Unabhängigkeit der Kontinuumshypothese über der Theorie ZFC: ZF kann das Auswahlaxiom weder beweisen noch widerlegen, ZFC kann die Kontinuumshypothese weder beweisen noch widerlegen (vorausgesetzt ZF – und damit ZFC – ist widerspruchsfrei). Weiter kann man Modelle von ZFC konstruieren, die „ $V \neq L$ “ erfüllen, d. h. ZFC kann nicht beweisen und nicht widerlegen, daß alle Mengen konstruktibel sind.

Für die Resultate von Gödel und Cohen ist ein formaler logischer Rahmen unerlässlich, die Resultate sind präzise mathematische Sätze über die Grenzen der in der Prädikatenlogik erster Stufe formulierten Zermelo-Fraenkel-Mengenlehre. Zwischen den Resultaten von Gödel und Cohen liegt mehr als ein Vierteljahrhundert, davor liegen große Anstrengungen, die von Zermelo informal begründete Axiomatik zu erweitern und zu formalisieren. Und insgesamt hat es 85 Jahre gedauert, bis man zeigen konnte, daß die von Cantor 1878 vermutete Kontinuumshypothese selbst mit den relativ starken Mitteln, die die Theorie ZFC zur Verfügung stellt, nicht zu lösen ist.

Einige Themen der modernen Mengenlehre

Einige große Gebiete der Mengenlehre am Beginn des 21. Jahrhunderts sind: Große Kardinalzahlen, innere Modelle, Forcing, deskriptive Mengenlehre/Axiom der Determiniertheit, unendliche Kombinatorik, Untersuchung von ZFC. Die seit Hausdorff und Mahlo untersuchten *großen Kardinalzahlen* bilden ein Zentrum des Interesses. Axiome, die die Existenz dieser Zahlen fordern, erscheinen heute als die natürlichsten Erweiterungen von ZFC. Ein *inneres Modell* ist eine transitive Klasse, innerhalb derer alle ZFC-Axiome gelten und die alle Ordinalzahlen enthält. Gödels inneres Modell L kann nur relativ kleine große Kardinalzahlen in sich tragen, erst die von Ronald Jensen initiierte Theorie und dann u. a. von John Steel weiterentwickelte Theorie der Kernmodelle liefert kanonische innere Modelle für große Kardinalzahlen. Die Theorie des *Forcing* ist, insbesondere durch Saharon Shelah, zur Reife gebracht worden, und auch hier bilden die großen Kardinalzahlen einen Schwerpunkt. Die *deskriptive Mengenlehre* baut auf den Vorarbeiten von Cantor, Hausdorff, Alexandrov, Lusin und Suslin auf, und untersucht definierbare Mengen von reellen Zahlen. Eine große Rolle spielt hier das sog. Axiom der Determiniertheit (AD) von Jan Mycielski (geb. 7.2.1932) und Hugo Steinhaus (14.1.1887 – 25.2.1972) aus dem Jahr 1962, ein Axiom über die Existenz von Gewinnstrategien für unendliche

Zweipersonenspiele, das in voller Stärke dem Auswahlaxiom widerspricht, aber in seinen schwachen Versionen Regularitätseigenschaften definierbarer Mengen sicherstellt. Auch hier wiederum sind die Beziehungen zu großen Kardinalzahlen sehr eng, und (AD) ist eng verwoben mit den sogenannten Woodin-Kardinalzahlen, benannt nach dem Mengentheoretiker Hugh Woodin. Die *unendliche Kombinatorik* versucht, die logische Stärke von kombinatorischen Prinzipien zu ermitteln. Dabei tritt das Äquikonsistenzphänomen auf, daß nämlich diese recht ungeordnet auftretenden Prinzipien jeweils gleichstark zu großen Kardinalzahlaxiomen sind. Letztere sind linear geordnet, und so erhalten natürliche mengentheoretische Aussagen eine unerwartete lineare Struktur. Insgesamt zeigen die Untersuchungen im Reich der über ZFC hinausgehenden Prinzipien, daß das Chaos dort nicht, wie man befürchten könnte, die Oberhand behält. Es gibt daneben nicht zuletzt auch Untersuchungen, die ganz in ZFC verbleiben: Es gibt in ZFC z. B. tiefliegende kombinatorische Ergebnisse und Resultate über Kardinalzahlarithmetik. Eine umfassende Theorie zur Kardinalzahlarithmetik in ZFC wurde in den 90er Jahren von Shelah entwickelt.

Ob sich die Frage nach der Kardinalität des Kontinuums irgendwann in überzeugender Weise doch beantworten läßt, ist weiter offen. Die Resultate von Gödel und Cohen zeigen ja nur, daß man dies nicht in ZFC tun kann. In letzter Zeit erschienen Beiträge zum Kontinuumsproblem von Hugh Woodin, die darauf hindeuten könnten, daß die Kontinuums-hypothese mit besseren Gründen als falsch denn als wahr anzusehen ist. Mehr wird die Zukunft zeigen.

Woodin (2001): „So, is the Continuum Hypothesis solvable? Perhaps I am not completely confident the ‘solution’ I have sketched is the solution, but it is for me convincing evidence that there *is* a solution. Thus, I now believe the Continuum Hypothesis is solvable, which is a fundamental change in my view of set theory. While most would agree that a clear resolution of the Continuum Hypothesis would be a remarkable event, it seems relatively few believe that such a resolution will ever happen.

Of course, for the dedicated skeptic there is always the ‘widget possibility’. This is the future where it is discovered that instead of sets we should be studying widgets. Further, it is realized that the axioms for widgets are obvious and, moreover, that these axioms resolve the Continuum Hypothesis (and everything else). For the eternal skeptic, these widgets are the integers (and the Continuum Hypothesis is resolved as being meaningless)...

The view that progress towards resolving the Continuum Hypothesis must come with progress on resolving all instances of the Generalized Continuum Hypothesis seems too strong. The understanding of $H(\omega) [= \{x \mid \text{für alle } x_n \in x_{n-1} \in \dots \in x_1 \in x_0 = x \text{ gilt: } x_i \text{ ist endlich für alle } 0 \leq i \leq n\} = \{x \mid x \text{ ist erblich endlich}\}]$ did not come in concert with an understanding of $H(\omega_1) [= \{x \mid \text{für alle } x_n \in x_{n-1} \in \dots \in x_1 \in x_0 = x \text{ gilt: } x_i \text{ ist abzählbar für alle } 0 \leq i \leq n\} = \{x \mid x \text{ ist erblich abzählbar}\}]$, and the understanding of $H(\omega_1)$ failed to resolve even the basic mysteries of $H(\omega_2) [= \{x \mid x \text{ ist erblich von der Kardinalität } \leq \aleph_1\}]$. The universe of sets is a large place. We have just barely begun to understand it.“

5. Literatur

Inhalt :

1. Die mathematischen Schriften von Georg Cantor
2. Die mathematischen Schriften von Felix Hausdorff
3. Die mathematischen Schriften von Ernst Zermelo
4. Mathematische Schriften anderer Autoren
5. Gesammelte Werke, Briefe und Aufsatzsammlungen
6. Ältere Bücher zur Mengenlehre
7. Bücher zur Mengenlehre bis etwa 1980
8. Neuere Bücher zur Mengenlehre
9. Bücher zu speziellen Themen der Mengenlehre
10. Bücher zur mathematischen Logik
11. Historische Arbeiten
12. Philosophische Schriften und Anthologien
13. Nichtmathematische Schriften von Georg Cantor
14. Schriften von Paul Mongré

Das folgende Literaturverzeichnis mit etwa 500 Einträgen bietet nicht nur „Nachweise und Weiterführendes“. Es ist auch zur „Lektüre“ gedacht, nicht nur zum Nachschlagen. Demgemäß ist es strukturiert und behandelt wissenschaftliche, biographische, allgemein-historische und philosophische Komplexe separat. Wir beschreiben im folgenden die einzelnen Teile des Verzeichnisses.

Abschnitte 1 – 3 :

Mathematische Schriften von Cantor, Hausdorff und Zermelo

Die Abschnitte beinhalten weitgehend vollständige Zusammenstellungen. Sie spiegeln den wissenschaftlichen Lebenslauf der drei Hauptpersonen wider, und unterstützen so die Biographien im Hauptteil des Buches. (Das Verzeichnis von Hausdorff folgt [Hausdorff 2001f], das von Zermelo einer Internet-Zusammenstellung von Volker Peckhaus.) Wir hoffen, daß der Leser einige Details bemerkenswert finden wird, die sich aus den Listen herausfiltern lassen: Etwa die wissenschaftliche Herkunft von Hausdorff und Zermelo aus anwendungsorientierten Gebieten, der Astronomie bei Hausdorff und der Variationsrechnung bei Zermelo. Daß Zermelo über das Zerbrechen eines Zuckerstücks nachgedacht hat, ist ja vielleicht bei dieser konfliktfähigen Figur nicht ohne Witz. Man sieht, daß Zermelo, im Gegensatz zu Hausdorff, die praktische Seite der Mathematik nie verlassen hat. Bei Hausdorff kann man beobachten, daß er in seiner Forschung zu bearbeiteten Fragen und ganzen Gebieten nach einer gewissen produktiven Phase nicht mehr zurückkehrt. Fruchtbare Jahre oder Schaffenskrisen sowie kritische oder entspannte Lebenssituationen aus den Listen zu diagnostizieren ist spekulativ, aber ein Vergleich mit den Biographien zeigt, daß man oft richtig liegt.

Wer Lust auf die Lektüre von wissenschaftlichen Artikeln der drei Mathematiker hat, dem sei hier empfohlen: Cantors sechsteilige Serie über lineare Punktmannigfaltigkeiten [Cantor 1879b, 1880d, 1882b, 1883a, 1883b, 1884b], und seine zweiteiligen „Beiträge“ [Cantor 1895a, 1897]. Diese Artikel finden sich auch in [Zermelo 1932]. Von Hausdorff: Seine große Arbeit über lineare Ordnungen [Hausdorff 1908]. Von Zermelo: Sein erster und zweiter Beweis des Wohlordnungssatzes [Zermelo 1904a, 1908a] und seine Axiomatisierung [Zermelo 1908b]. Daneben ist [Zermelo 1930] ein fesselnder Artikel.

Abschnitt 4 : Mathematische Schriften anderer Autoren

Die frühen Arbeiten zur Mengenlehre bis etwa 1930 stehen hier im Vordergrund, bei gelegentlichen Vorausgriffen auf neuere Literatur, die der Text des Buches nötig macht. Der Leser wird sehen, daß dieses Buch bei weitem nicht alles abdeckt, was es an Interessantem aus der frühen Zeit der Mengenlehre zu berichten gäbe.

Hier kann als Empfehlung nur eine kleine Auswahl von Titeln genannt werden, die sicher nicht repräsentativ für diesen Abschnitt ist: Die Bücher von Dedekind [Dedekind 1872, 1888] und Frege [Frege, 1884], allesamt vielfach neu aufgelegt. Dann für die Mengenlehre die Artikel [Bernstein 1905], [Hartogs 1915], [von Neumann 1923, 1928], [Skolem 1923]. Für die Logik allgemein die Arbeiten von Kurt Gödel. Speziell für die frühe deskriptive Mengenlehre die Arbeiten von Nikolai Lusin. In Richtung große Kardinalzahlen die Arbeiten von Paul Mahlo, und einige Arbeiten von Tarski.

Abschnitt 5 : Gesammelte Werke, Briefe und Aufsatzsammlungen

Der Abschnitt gibt einen Überblick über gesammelte Werke. Während eine Gesamtausgabe für Hausdorff in jüngster Zeit in Angriff genommen wurde, steht eine Neuauflage der Schriften von Cantor noch aus, und von Zermelo fehlt eine Werkausgabe bislang völlig.

Neben der neuen Hausdorff-Edition [Hausdorff 2001f] ist die von Zermelo besorgte Sammlung der Werke Cantors [Cantor 1932] und die Ausgabe der Briefe [Cantor 1991] besonders empfehlenswert. Die sehr gut gemachten Bücher [Felgner 1979] und [van Heijenoort 1967] versammeln wichtige Originalartikel. Gödel ist immer interessant, und den an der mathematischen Logik interessierten Leser wird hier die Werkausgabe [Gödel 1986ff] besonders erfreuen. Auch die dreibändige Werkausgabe der Schriften von Dedekind [Dedekind 1930 – 1932] ist eine Großtat.

Abschnitte 6 – 10 : Mengentheoretische Lehrbuchliteratur

Auch hier geht es nicht nur um die Angabe von Referenzen: In der Buchliteratur spiegelt sich der Zustand einer mathematischen Theorie. In der Regel häufen sich neue Bücher einige Jahre nach besonders aufregenden Phasen der Forschung. Derartige Häufungen wichtiger mengentheoretischer Lehrbuchliteratur bilden sich etwa um die Jahre 1910, 1975 und 1995.

Die einführenden, elementaren Lehrbücher zur Mengenlehre lassen sich relativ natürlich in drei zeitliche Blöcke einteilen, die die Zeitspannen von etwa 1900 – 1950, 1950 – 1980 sowie 1980 bis heute umfassen.

Abschnitt 6 : Ältere Bücher zur Mengenlehre

Wer bei dem in diesem Buch dargestellten Stoff noch verweilen und die originale Luft atmen möchte, findet in dieser Liste alles, was die Geschichte hinterlassen hat, und fast al-

les, was das Herz begehrt. Die ersten fünfzig Jahre der Mengenlehre, die in diesen Büchern zeitnah dargestellt werden, haben im Kleinen unendlich viele fesselnde Feinheiten zu bieten, und sind als Ganzes geistes- und kulturgeschichtlich unerreicht.

Die Empfehlungen hier sind eindeutig: [Hausdorff 1914] und [Fraenkel 1928]. Die „Grundzüge“ liegen mittlerweile in einer Neuausgabe samt umfangreichem Apparat vor [Hausdorff 2002]. Die zweite Auflage von Hausdorffs Buch von 1914 [Hausdorff 1927] enthält weniger und mehr als die erste: Weniger Ordnungstheorie, weniger Topologie, weniger stilistisch ausgereifte Kommentare, dafür aber mehr deskriptive Mengenlehre. Insgesamt bleibt die Empfehlung bei [Hausdorff 1914].

Die Bücher von Schoenflies und Hessenberg tragen das Prädikat „historisch besonders wertvoll“. Für die deskriptive Mengenlehre ist [Lusin 1930] ein Klassiker.

Abschnitt 7 : Bücher zur Mengenlehre bis etwa 1980

Der zweite – nicht vollständige – Block zur Buchliteratur spiegelt die Zeit der weiten Verbreitung der Mengenlehre, die nach dem Krieg einsetzte. Die Mathematik wurde auf die Sprache der Mengenlehre umgestellt und axiomatisch entwickelt, insbesondere durch die Arbeiten der französischen Bourbaki-Schule. Im Gegensatz zu topologischen Begriffen gelangten aber zentrale inhaltliche Konzepte der allgemeinen Mengenlehre wie Ordinal- und Kardinalzahlen nicht in den Kanon des Grundwissens der Mathematik, und die Theorie dümpelte ein wenig vor sich hin. Die Erfindung des Forcing durch Paul Cohen 1963 brachte die Mengenlehre dann erneut in die Schlagzeilen, und weckte sie aus ihrem Schönheitsschlaf. Die neue, sehr flexible Methode zog viele junge Forscher an, und viele klassische Probleme konnten beinahe im Wochentakt gelöst werden. Die Mengenlehre entwickelte sich zu einer vielfältigen, blühenden Theorie, in der heute das die Dinge erneut in Bewegung setzende Forcing nur ein Stichwort unter vielen anderen ist.

In den siebziger Jahren sickerten Konzepte der Mengenlehre sogar bis zu den Grundschulen hinab, wurden dann aber schnell von ungezählten Eltern wieder vertrieben, die bei den Hausaufgaben ihrer Kinder an ihrem Bildungsvorsprung zu zweifeln begannen. An den Universitäten in Deutschland erholte sich die durch den Krieg und die Vertreibung von Wissenschaftlern völlig unterbrochene Tradition nur langsam. Eine für die Mengenlehre als Ganzes und für die heutige mengentheoretische Forschung in Deutschland zentrale Figur ist sicherlich Ronald Jensen.

Nur wenige einführende Bücher der Liste sind noch als Ganzes genießbar; der frische Wind der ersten Zeit ist vorbei, und der Staub, den die Geschichte hinterlassen hat, bleibt oft liegen. Die Ausnahme bildet das Buch [Halmos 1960], das viel zur Verbreitung einiger elementarer Ideen der Mengenlehre beigetragen hat, auch unter Laien. Es ist auch heute noch lesenswert, und es ist leicht zugänglich. Auf höherem Niveau ragen die Bücher von Kuratowski und Mostowski heraus.

Besonders interessante Mitglieder der Liste sind die Titel, die sich mit Forcing auseinandersetzen; [Cohen 1966] ist ein Klassiker, und daneben bietet [Jensen 1967] eine frühe Darstellung der Modellkonstruktion durch die Erzwingungsmethode.

Abschnitt 8 : Neuere Bücher zur Mengenlehre

Der dritte Teil der Literatur zur Mengenlehre schließlich zeigt Alternativen und Ergänzungen zum vorliegenden Buch und verweist auf Texte, die weit über das hier behandelte Material hinausgehen.

Hinsichtlich einführender Literatur ist vieles Geschmackssache. Wir raten dem Leser, der eine Ergänzung sucht, sich unter [Ebbinghaus 2003], [Friedrichsdorf/ Prestel 1985], [Hrbacek/ Jech 1999], [Oberschelp 1994] sowie [Moschovakis 1994] den ihn am meisten ansprechenden Text auszuwählen.

Die herausragenden Lehrbücher ab 1980 für die höhere Mengenlehre sind zweifellos [Jech 1978/2003] und [Kunen 1980]. Das sehr kompakt geschriebene Buch von Jech war schon bei seinem ersten Erscheinen eine Enzyklopädie, und die erweiterte Neufassung von 2003 bietet dem Leser einen einzigartigen Überblick über den Stand der Dinge: [Jech 2003] ist in vielen Punkten *die* Empfehlung zur modernen Mengenlehre. Ein interessantes Buch, in dem man viele Details finden kann, die anderswo verloren gegangen sind, ist [Levy 1979].

[Deiser 2007] behandelt die reellen Zahlen unter verschiedenen Gesichtspunkten, insbesondere solchen der deskriptiven Mengenlehre und der Grundfragen der Maßtheorie.

Abschnitt 9 : Bücher zu speziellen Themen der Mengenlehre

Ab etwa 1990 ist eine ungewöhnliche Vielzahl von meistens sehr fortgeschrittenen und spezialisierten Büchern erschienen. Manche von ihnen sind auch für Experten schwer zugänglich und eher Forschungsberichte als Lehrbücher. Die Liste zeigt die Lebendigkeit und den Reichtum in der Forschung der letzten drei Jahrzehnte. Inhalte und Stil der Bücher zeigen aber auch, daß sich das Cantorsche Paradies in einen Hightech-Maschinenpark verwandelt hat, mit allen damit verbundenen Vor- und Nachteilen.

Zu den freundlicheren Titeln: [Jech 1973], [Rubin 1963, 1985] und [Herrlich 2006] sind Monographien über das Auswahlaxiom. Das Buch [Kanamori 1994] gibt eine umfassende, gut lesbare und historisch orientierte Einführung in die Theorie der großen Kardinalzahlen. [Kechris 1995] ist ein für jeden Mathematiker geschriebener Text zur deskriptiven Mengenlehre, der auch Analytiker begeistern kann. Kenntnisse in mathematischer Logik sind nicht erforderlich. [Moschovakis 1980] ist dagegen der Klassiker zur logisch formulierten deskriptiven Mengenlehre. Das Buch [Devlin 1984] widmet sich ganz dem konstruktiblen Universum von Gödel und seiner eingehenden Untersuchung durch Ronald Jensen. [Martin 200?] ist ein gerade entstehendes, umfassendes Buch über ein spezielles Axiom, das Axiom der Determiniertheit, welches in voller Stärke dem Auswahlaxiom widerspricht, in abgeschwächten Formen aber garantiert, daß sich definierbare Teilmengen von $\mathbb{R}$ ordentlich benehmen.

Zwei Bücher mit Anwendungen: [Dales / Woodin 1987] diskutiert eine Anwendung der Erzwingungsmethode in der Theorie der Banachalgebren, und beinhaltet eine selbstständige Darstellung der benötigten Werkzeuge aus beiden Bereichen. [Dehornoy 2000] ist eine glänzende Monographie über Zöpfe und eine skurrile algebraische Identität. Der Mengenlehre kommt hier überraschenderweise eine problemgeschichtlich und inhaltlich wichtige Rolle zu, die der letzte Teil des Buches deutlich macht.

Schließlich zu einigen hochspezialisierten Texten. Die Bücher von Tony Dodd, John Steel, William Mitchell und das auf Manuskripten von Ronald Jensen aufbauende Buch von Zeman behandeln die in erster Linie durch die Autoren selbst erforschte Kernmodelltheorie, eine Weiterentwicklung der Gödelschen Theorie des konstruktiblen Universums in Richtung große Kardinalzahlaxiome. Die Bücher von Shelah behandeln die Resultate des Autors zur Kardinalzahlarithmetik in ZFC bzw. spezielle Techniken der iterierten Erzwingungsmethode. Das große Buch von Hugh Woodin [Woodin 1999] ist ein Meilenstein in jeder Hinsicht, geschrieben von einem überragenden Mengentheoretiker. (Jensen, Shelah, Steel und Woodin bilden wahrscheinlich das mengentheoretische Viergestirn der heutigen Zeit. Ein Buch über sie zu schreiben, wie man es über Cantor, Hausdorff und Zermelo tun kann, wird aber wohl unmöglich sein: Alleine ein vollständiges Literaturverzeichnis der Arbeiten von Shelah hätte über 800 Einträge, und das erste Tausend wird wohl voll werden.)

Abschnitt 10 : Bücher zur mathematischen Logik

Für die höhere Mengenlehre ist etwas allgemeine Logik und Modelltheorie ebenso unvermeidlich wie die Maßtheorie für die Wahrscheinlichkeitstheorie.

Empfehlenswerte Einführungen in die Prädikatenlogik sind [Ebbinghaus / Flum / Thomas 1996], [Mates 1997], [Prestel 1986] und [Rautenberg 1996]; das letztgenannte enthält eine sehr gute Darstellung der Gödelschen Unvollständigkeitssätze. Schon etwas ältere Klassiker der Liste sind [Shoenfield 1967] und [Tarski 1977] aus dem englischen Sprachraum, und [Hermes 1976] auf Deutsch.

In der Logik gibt es neben der Mengenlehre noch die Berechenbarkeitstheorie, die Beweistheorie und die Modelltheorie, und zu diesen Themenfeldern sind einige Lehrbücher aufgeführt. Der Klassiker zur Modelltheorie ist [Chang / Keisler 1990], der zur Berechenbarkeitstheorie [Rogers 1987], der zur Beweistheorie [Schütte 1960]. Letzterer ist für Anfänger schwer zugänglich. Eine neuere Einführung in die Beweistheorie ist [Schwichtenberg / Troelstra 2000].

Wer nur etwas Rüstzeug für weitere Studien in der Mengenlehre sucht, ist mit einem einführenden Text, der etwas mehr Modelltheorie behandelt, gut bedient. Der Leser möge sich hierbei nicht durch das etwas ungewöhnliche Verhältnis zwischen Logik und Mengenlehre verwirren lassen. Prädikatenlogik erster Stufe und Modelltheorie kann man gut innerhalb der Mengenlehre betreiben. Die formale Mengenlehre ist eine mit finiten Mitteln formulierte Theorie, wie im dritten Abschnitt beschrieben. Und daß in dieser Theorie dann nicht nur Zahlentheorie und Algebra, sondern auch eine allgemeine Theorie über Syntax und Semantik entwickelt werden kann, in der sich ihre formale Natur spiegelt, ist keineswegs zirkulär. Vom Standpunkt der Mengenlehre aus ist die Prädikatenlogik eine mathematische Disziplin wie andere auch. Andererseits muß man als Logiker keinerlei mengentheoretische Orthodoxie unterschreiben. Die Mengenlehre wird dann eher als Sprache verwendet denn als Theorie über abstrakte unendliche Objekte. Viele Fragen der Logik sind rein finiter Natur, und damit sehr konkret und „computerisierbar“. Kurz: Logik ist bereits ohne aktuell unendliche Objekte von großem Interesse. In jedem Fall ist es didaktisch einfacher, die Logik innerhalb einer naiven Mengenlehre zu entwickeln, und die meisten Lehrbücher gehen diesen klaren und durch die Geschichte freigeräumten Weg.

Abschnitt 11 : Historische Arbeiten

Der Abschnitt gibt eine Auswahl an historischen Artikeln und einigen Büchern zur Geschichte der Mengenlehre, Logik und mathematischer Grundlagenforschung. Viele Details der frühen Geschichte sind in den letzten 25 Jahren untersucht worden, und zur Person Cantor liegen verschiedene Biographien vor, die wissenschaftliche, philosophische und menschliche Aspekte diskutieren.

Die (Standard-)Empfehlungen für Bücher, die sich auf die Person Cantor und seine Theorie konzentrieren, wären [Dauben 1979], [Hallett 1984], [Meschkowski 1967] sowie [Purkert / Illgauds 1987]. Zwei Bücher mit einem weiter gefaßten Rahmen sind [Ferreirós 1999] und [Grattan-Guinness 2000]. Ferreirós holt insbesondere Dedekind und Riemann mit ins Boot, und im Buch von Grattan-Guinness findet man viel über Bertrand Russell. Auf Band II der neuen Hausdorff-Ausgabe [Hausdorff 2002] wurde bereits hingewiesen; er erfreut nicht nur durch die Faksimile-Wiedergabe der „Grundzüge“, sondern auch durch begleitende Essays verschiedener Autoren. Ein interessanter Sammelband zu Hausdorff ist [Eichhorn / Thiele 1994]. Band I der Hausdorff-Ausgabe wird eine Biographie enthalten. Unbedingt ist das Buch von Moore (1982) über Zermelos Auswahlaxiom zu nennen, das auch viele Informationen über die Entwicklung der Mengenlehre zwischen Zermelo und Cohen enthält (etwa über die polnische Schule um Sierpiński). Die Zermelo-Biographie [Ebbinghaus 2007] ist ganz hervorragend. Eine von Ebbinghaus u. a. betreute Werkausgabe von Zermelo ist in Arbeit. Ein Klassiker zur Geschichte der Mengenlehre ist die Artikelserie [Jourdain 1906 – 1914]. Schließlich sei hier auch noch die Sammlung [Hintikka 1995] genannt.

Abschnitt 12 : Philosophische Schriften und Anthologien

Eine Auswahl. Ganz naiv formuliert ist die große Frage die nach der Wahrheit, und was abstrakte Existenz ist. Kann man zum Phänomen, daß das Universum des Unendlichen ähnlich reich ausgestattet ist wie der wirkliche Kosmos, noch irgendetwas Sinnvolles sagen, oder muß man das staunend so hinnehmen? Das Staunen steht am Anfang der Philosophie, und danach wird es schwierig und zuweilen auch schwammig. Gutes Kartenmaterial für diesen Dschungel zu bekommen, ist nicht leicht. Gegen den Blick trübende Schlangenbisse hilft die Rückkehr zum kristallklaren Definition-Satz-Beweis-Schema der Mathematik.

Die erste Empfehlung ist Platon, der gar nicht aufgelistet ist. Überall leicht zu bekommen. Wer etwas Neuere zu Platon lesen will, wird vielleicht an [Moravcsik 1992] seine Freude haben.

Es gibt eine Reihe von neueren Anthologien, mit denen der Leser den ihn interessierenden Weg beginnen kann: [Dales / Olivieri 1998], [Hart 1996], [Jacquette 2002], [Link 2004], [Schirn 1998], [Shanker 2001], [Thiel 1982], [Tymoczko 1998].

Den Anthologien stehen etwa die Monographien und Essay-Sammlungen [Balaguer 1998], [Bernays 1976], [George / Velleman 2002], [Maddy 1990, 1997], [Parsons 2005], [Thiel 1989], [Ties 1989] gegenüber. [Thiel 1995] ist als eine Einführung empfehlenswert. [Beth 1968] ist ein Beispiel für ein Werk, das auf den ersten Blick viel verspricht, aber den Autor dieses Buches enttäuscht hat.

Auch viele Bücher zur mathematischen Logik enthalten oft hervorragende Passagen, die sich mit philosophischen Themen und den zugehörigen historischen Ereignissen befassen. Ein Beispiel ist das in Abschnitt 10 genannte Buch [Kleene 1968] und seine Diskussion des Grundlagenstreits zwischen Hilbert und Brouwer.

Abschnitt 13 : Nichtmathematische Schriften von Georg Cantor

Cantor hat in den 80er Jahren des 19. Jahrhunderts über die von Delia Bacon vorgebrachte These „Shakespeare=Bacon“ geschrieben und geforscht. Der gleiche Nachname „Bacon“ ist rein zufällig, zunächst ohne Bedeutung, später wurde er aber von Frau Bacon in ihre Theorie integriert. Die Bacon-These war damals relativ populär, heute ist es eher der Earl of Oxford, der Shakespeare gewesen sein soll. (Nicht nur als Mathematiker vertraut der Autor des vorliegenden Buches auf Shakespeare=Shakespeare.)

Cantor hat versucht, Beweise für die These zu finden, und hat einiges zum Thema veröffentlicht. Eine genauere Darstellung findet sich in [Purkert / Ilgauds 1987].

Weiter hat Cantor religiös-mystische Schriften herausgegeben und verfaßt. Für eine Neuausgabe und Kommentierung dieser Schriften muß man wohl auf eine neue Gesamtausgabe warten, die dann die alte Zermelosche Ausgabe von 1932 ablösen wird.

Abschnitt 14 : Schriften von Paul Mongré

Hier schließlich die Schriften von Paul Mongré (zitiert nach [Hausdorff 2001f]). Einige Bemerkungen zu dieser Gestalt findet der Leser in der Hausdorff-Biographie. Die Bände VI und VII der Hausdorff-Edition werden diese Schriften von Mongré wieder besser zugänglich machen. Bis dahin lesen wir Hausdorff.

1. Die mathematischen Schriften von Georg Cantor

- 1867** De aequationibus secundi gradus indeterminatis / Georgius Cantor;
Dissertation (Diss. phil.) mit einer Vita von Georg Cantor: Schultz, Berlin.
- 1868** Zwei Sätze aus der Theorie der binären quadratischen Formen; *Zeitschrift für Mathematik und Physik* 13 (1868), S. 259 – 261. (Innerhalb der Rubrik: *Kleinere Mitteilungen.*)
- 1869a** Über die einfachen Zahlensysteme; *Zeitschrift für Mathematik und Physik* 14 (1869), S. 121 – 128.
- 1869b** Zwei Sätze über eine gewisse Zerlegung der Zahlen in unendliche Produkte; *Zeitschrift für Mathematik und Physik* 14 (1869), S. 152 – 158. (Innerhalb der Rubrik: *Kleinere Mitteilungen.*)
- 1869c** De transformatione formarum ternariarum quadraticarum / Georgius Cantor;
Habilitationsschrift. Hendel, Halle.
- 1870a** Über einen die trigonometrischen Reihen betreffenden Lehrsatz; *Journal für die reine und angewandte Mathematik* 72 (1870), S. 130 – 138.
- 1870b** Beweis, daß eine für jeden reellen Wert von x durch eine trigonometrische Reihe gegebene Funktion $f(x)$ sich nur auf eine einzige Weise in dieser Form darstellen läßt; *Journal für die reine und angewandte Mathematik* 72 (1870), S. 139 – 142.
- 1871a** Notiz zu dem Aufsatz: Beweis, daß eine für jeden reellen Wert von x durch eine trigonometrische Reihe gegebene Funktion $f(x)$ sich nur auf eine einzige Weise in dieser Form darstellen läßt. Bd. 72, S. 139 dieses Journals; *Journal für die reine und angewandte Mathematik* 73 (1871), S. 294 – 296.
- 1871b** Über trigonometrische Reihen; *Mathematische Annalen* 4 (1871), S. 139 – 143.
- 1871c** Rezension von: H. Hankel, Untersuchungen über die unendlich oft oszillierenden und unstetigen Funktionen. Tübingen, Universitätsprogramm, 1870; *Literarisches Centralblatt* 7 (18. 2. 1871), S. 150 – 151.
- 1872a** Über die Ausdehnung eines Satzes aus der Theorie der trigonometrischen Reihen; *Mathematische Annalen* 5 (1872), S. 123 – 132.
- 1872b** Algebraische Notiz; *Mathematische Annalen* 5 (1872), S. 133 – 134.
- 1873** Historische Notizen über die Wahrscheinlichkeitsrechnung; *Bericht über die Sitzungen der Naturforschenden Gesellschaft zu Halle im Jahre 1873* (1873), S. 34 – 42.
- 1874** Über eine Eigenschaft des Inbegriffes aller reellen algebraischen Zahlen; *Journal für die reine und angewandte Mathematik* 77 (1874), S. 258 – 262.
- 1878** Ein Beitrag zur Mannigfaltigkeitslehre; *Journal für die reine und angewandte Mathematik* 84 (1878), S. 242 – 258.
- 1879a** Über einen Satz aus der Theorie der stetigen Mannigfaltigkeiten; *Nachrichten von der Königl. Gesellschaft der Wissenschaften und der Georg-Augusts-Universität (Göttinger Nachrichten)* (1879), S. 127 – 135.
- 1879b** Über unendliche, lineare Punktmannigfaltigkeiten. 1; *Mathematische Annalen* 15 (1879), S. 1 – 7.
- 1880a** Bemerkung über trigonometrische Reihen; *Mathematische Annalen* 16 (1880), S. 113 – 114.

- 1880b** Fernere Bemerkung über trigonometrische Reihen; *Mathematische Annalen* 16 (1880), S. 267 – 269.
- 1880c** Zur Theorie der zahlentheoretischen Funktionen; *Nachrichten von der Königl. Gesellschaft der Wissenschaften und der Georg-Augusts-Universität (Göttinger Nachrichten)* (1880), S. 161 – 169.
Daneben auch in: *Mathematische Annalen* 16 (1880), S. 583 – 588.
- 1880d** Über unendliche, lineare Punktmannigfaltigkeiten. 2; *Mathematische Annalen* 17 (1880), S. 355 – 358.
- 1881** Rezension von: Briefwechsel zwischen Gauß und Bessel. Herausgegeben auf Veranlassung der Königl. preussischen Akademie der Wissenschaften. Engelmann, Leipzig, 1880; *Deutsche Literaturzeitung* 2 (1881), S. 1082.
- 1882a** Über ein neues und allgemeines Kondensationsprinzip der Singularitäten von Funktionen; *Mathematische Annalen* 19 (1882), S. 588 – 594.
- 1882b** Über unendliche, lineare Punktmannigfaltigkeiten. 3; *Mathematische Annalen* 20 (1882), S. 113 – 121.
- 1883a** Über unendliche, lineare Punktmannigfaltigkeiten. 4; *Mathematische Annalen* 21 (1883), S. 51 – 58.
- 1883b** Über unendliche, lineare Punktmannigfaltigkeiten. 5; *Mathematische Annalen* 21 (1883), S. 545 – 591.
(Oft zitiert als „S. 545 – 586“. Die Seiten 587 – 591 beinhalten Anmerkungen von Georg Cantor zu seinem Artikel.)
- 1883c** Grundlagen einer allgemeinen Mannigfaltigkeitslehre. Ein mathematisch-philosophischer Versuch in der Lehre des Unendlichen; Teubner, Leipzig.
(Um ein Vorwort ergänzte Ausgabe von 1883b.)
- 1883d** Sur divers théorèmes de la théorie des ensembles de points situés dans un espace continu à n dimensions. Première communication; *Acta Mathematica* 2 (1883), S. 409 – 414.
- 1884a** De la puissance des ensembles parfaits de points. Extrait d’une lettre adressée à l’éditeur; *Acta Mathematica* 4 (1884), S. 381 – 392.
- 1884b** Über unendliche, lineare Punktmannigfaltigkeiten. Nr. 6; *Mathematische Annalen* 23 (1884), S. 453 – 488.
- 1884c** Rezension von: H. Cohen, Das Prinzip der Infinitesimalmethode und seine Geschichte. Dümmler, Berlin, 1883; *Deutsche Literaturzeitung* 4 (1884), S. 266 – 268.
- 1884x** Prinzipien einer Theorie der Ordnungstypen. Erste Mitteilung; eine zu Lebzeiten Cantors nicht veröffentlichte Arbeit mit Datum 6.11.1884. Letztendlich veröffentlicht und mit Dokumenten versehen in: *Acta Mathematica* 124 (1970), S. 65 – 107. Herausgegeben von Ivor Grattan-Guinness.
- 1885a** Ludwig Scheeffer (1859 – 1885). Nekrolog; *Bibliotheca Mathematica* 2 (1885), Spalten 197 – 199.
- 1885b** Rezension von: Gottlob Frege, Die Grundlagen der Arithmetik. Koebner, Breslau, 1884; *Deutsche Literaturzeitung*, VI. Jahrgang (1885), S. 728 – 729.
- 1885c** Über die verschiedenen Ansichten in bezug auf die aktualunendlichen Zahlen; *Bibang till K. Svenska Vetenskaps-Akademiens Handlingar* 11, Nr. 19 (1885), S. 1 – 9.

- 1885d** Über verschiedene Theoreme aus der Theorie der Punktmengen in einem n -fach ausgedehnten stetigen Raume G_n . Zweite Mitteilung; *Acta Mathematica* 7 (1985), S. 105 – 124.
(Fortsetzung des französischen Artikels 1883d.)
- 1886a** Über die verschiedenen Standpunkte in Bezug auf das aktual Unendliche; *Zeitschrift für Philosophie und philosophische Kritik* 88 (1886), S. 224 – 233.
- 1886b** Über die verschiedenen Standpunkte in Bezug auf das aktual Unendliche. (Brief vom 4. 11. 1885 an G. Eneström, Redakteur der Bibliotheca mathematica in Stockholm); *Natur und Offenbarung. Organ zur Vermittlung zwischen Naturforschung und Glauben für Gebildete aller Stände* 32 (1886), S. 46 – 49.
- 1886c** Zum Problem des aktuellen Unendlichen; *Natur und Offenbarung. Organ zur Vermittlung zwischen Naturforschung und Glauben für Gebildete aller Stände* 32 (1886), S. 226 – 233.
- 1887** Mitteilungen zur Lehre vom Transfiniten. 1; *Zeitschrift für Philosophie und philosophische Kritik* 91 (1887), S. 81 – 125 und 252 – 270.
- 1888** Mitteilungen zur Lehre vom Transfiniten. 2; *Zeitschrift für Philosophie und philosophische Kritik* 92 (1888), S. 240 – 265.
- 1889** Bemerkung mit Bezug auf den Aufsatz: Zur Weierstraß'-Cantor'schen Theorie der Irrationalzahlen in Math. Annalen Bd. XXXIII, p. 154; *Mathematische Annalen* 33 (1889), S. 476. Die Bemerkung bezieht sich auf den Aufsatz von Eberhard Illigens in: *Mathematische Annalen* 33, S. 155 – 160. (Abweichend zu „p. 154“ im Titel.)
- 1890** Zur Lehre vom Transfiniten. Gesammelte Abhandlungen aus der Zeitschrift für Philosophie und philosophische Kritik von Georg Cantor, ord. Professor an der Universität Halle – Erste Abteilung; *Verlag von C. E. M. Pfeffer (Norbert Stricker), Halle*.
- 1892** Über eine elementare Frage der Mannigfaltigkeitslehre; *Jahresbericht der Deutschen Mathematiker-Vereinigung. Erster Band. 1890 – 91. 1* (1892), S. 75 – 78.
(Nachdruck bei Johnson, New York, 1960.)
- 1894** Vérification jusqu'à 1000 du théorème empirique de Goldbach; *Association Française pour l'Avancement des Sciences, C. R. de la 23^{me} Session (Caen 1894), Seconde Partie* (1895), S. 117 – 134.
- 1895a** Beiträge zur Begründung der transfiniten Mengenlehre. (Erster Artikel); *Mathematische Annalen* 46 (1895), S. 481 – 512.
- 1895b** Sui numeri transfiniti. Estratto d'una lettera di Georg Cantor a G. Vivanti – 13. Dec. 1893; *Rivista di Matematica* 5 (1895), S. 104 – 108.
(In deutscher Sprache. Briefdatum 13. 12. 1883)
- 1895c** Lettera di Georg Cantor a Giuseppe Peano; *Rivista di Matematica* 5 (1895), S. 108 – 109. (In deutscher Sprache. Briefdatum 20. Juli 1895)
- 1897** Beiträge zur Begründung der transfiniten Mengenlehre. (Zweiter Artikel); *Mathematische Annalen* 49 (1897), S. 207 – 246.
- 1903** Bemerkungen zur Mengenlehre; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 12 (1903), S. 519.
(Keine Publikation, sondern Erwähnung eines Vortrags von Cantor mit dem Titel „Bemerkungen zur Mengenlehre“.)
- 1905** Brief von Carl Weierstraß über das Dreikörperproblem; *Rendiconti del Circolo Matematico di Palermo* 19 (1905), S. 305 – 308.

2. Die mathematischen Schriften von Felix Hausdorff

- 1891** Zur Theorie der astronomischen Strahlenbrechung; *Dissertation. Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 43 (1891), S. 481 – 566.
- 1893 a** Zur Theorie der astronomischen Strahlenbrechung II; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 45 (1893), S. 120 – 162.
- 1893 b** Zur Theorie der astronomischen Strahlenbrechung III; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 45 (1893), S. 758 – 804.
- 1895** Über die Absorption des Lichtes in der Atmosphäre; *Habilitationsschrift. Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 47 (1895), S. 401 – 482.
- 1896** Infinitesimale Abbildungen in der Optik; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 48 (1896), S. 79 – 130.
- 1897** Das Risiko bei Zufallsspielen; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 49 (1897), S. 497 – 548.
- 1899** Analytische Beiträge zur nichteuklidischen Geometrie; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 51 (1899), S. 161 – 214.
- 1900** Zur Theorie der Systeme komplexer Zahlen; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 52 (1900), S. 43 – 61.
- 1901 a** Beiträge zur Wahrscheinlichkeitsrechnung; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 53 (1901), S. 152 – 178.
- 1901 b** Über eine gewisse Art geordneter Mengen; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 53 (1901), S. 460 – 475.
- 1903** Das Raumproblem; *Ostwalds Annalen der Naturphilosophie* 3 (1903), S. 1 – 23. (Abdruck der Antrittsvorlesung an der Universität Leipzig vom 4.7.1903.)
- 1904** Der Potenzbegriff in der Mengenlehre; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 13 (1904), S. 569 – 571.
- 1905** Bertrand Russell, The principles of mathematics; *Rezension. Vierteljahresschrift für wissenschaftliche Philosophie und Sociologie* 29 (1905), S. 119 – 124.
- 1906 a** Die symbolische Exponentialformel in der Gruppentheorie; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 58 (1906), S. 19 – 48.
- 1906 b** Untersuchungen über Ordnungstypen I, II, III; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 58 (1906), S. 106 – 169.

- 1907a** Untersuchungen über Ordnungstypen IV, V; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 59 (1907), S. 84 – 159.
- 1907c** Über dichte Ordnungstypen; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 16 (1907), S. 541 – 546.
- 1908** Grundzüge einer Theorie der geordneten Mengen; *Mathematische Annalen* 65 (1908), S. 435 – 505.
- 1909a** Die Graduierung nach dem Endverlauf; *Abhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 31 (1909), S. 295 – 334.
- 1909b** Zur Hilbertschen Lösung des Warringschen Problems; *Mathematische Annalen* 67 (1909), S. 301 – 305.
- 1914a** Grundzüge der Mengenlehre; *Veit & Comp., Leipzig. (Nachdrucke bei Chelsea, New York 1949, 1965, 1978.)*
- 1914b** Bemerkungen über den Inhalt von Punktmengen; *Mathematische Annalen* 75 (1914), S. 428 – 433.
- 1916** Die Mächtigkeit der Borelschen Mengen; *Mathematische Annalen* 77 (1916), S. 430 – 437.
- 1919a** Dimension und äußeres Maß; *Mathematische Annalen* 79 (1919), S. 157 – 179.
- 1919b** Der Wertvorrat einer Bilinearform; *Mathematische Zeitschrift* 3 (1919), S. 314 – 316.
- 1919c** Zur Verteilung der fortsetzbaren Potenzreihen; *Mathematische Zeitschrift* 4 (1919), S. 98 – 103.
- 1919d** Über halbstetige Funktionen und deren Verallgemeinerung; *Mathematische Zeitschrift* 5 (1919), S. 292 – 309.
- 1921a** Summationsmethoden und Momentfolgen. I; *Mathematische Zeitschrift* 9 (1921), S. 74 – 109.
- 1921b** Summationsmethoden und Momentfolgen. II; *Mathematische Zeitschrift* 9 (1921), S. 280 – 299.
- 1923a** Eine Ausdehnung des Parsevalschen Satzes über Fourierreihen; *Mathematische Zeitschrift* 16 (1923), S. 163 – 169.
- 1923b** Momentprobleme für ein endliches Intervall; *Mathematische Zeitschrift* 16 (1923), S. 220 – 248.
- 1924** Die Mengen G_8 in vollständigen Räumen; *Fundamenta Mathematicae* 6 (1924), S. 146 – 148.
- 1925** Zum Hölderschen Satze über $\Gamma(x)$; *Mathematische Annalen* 94 (1925), S. 244 – 247.
- 1927a** Mengenlehre; *zweite, stark umgearbeitete Auflage von „Grundzüge der Mengenlehre“, 1914. Walter de Gruyter, Berlin.*
- 1927b** Beweis eines Satzes von Arzelà; *Mathematische Zeitschrift* 26 (1927), S. 135 – 137.
- 1927c** Lipschitzsche Zahlensysteme und Studysche Nablafunktionen; *Journal für die reine und angewandte Mathematik* 158 (1927), S. 113 – 127.
- 1930a** Die Äquivalenz der Hölderschen und Cesàroschen Grenzwerte negativer Ordnung; *Mathematische Zeitschrift* 31 (1930), S. 186 – 196.
- 1930b** Erweiterung einer Homöomorphie; *Fundamenta Mathematicae* 16 (1930), S. 353 – 360.

- 1931/1932** Zur Theorie der linearen metrischen Räume; *Journal für die reine und angewandte Mathematik* 167 (1931/1932), S. 294 – 311.
- 1933a** Zur Projektivität der δ s-Funktionen; *Fundamenta Mathematicae* 20 (1933), S. 100 – 104.
- 1933b** Problem 58; *Fundamenta Mathematicae* 20 (1933), S. 286.
- 1934** Über innere Abbildungen; *Fundamenta Mathematicae* 23 (1934), S. 279 – 291.
- 1935a** Mengenlehre; *dritte Auflage*. Walter de Gruyter, Berlin. (Nachdruck bei Dover, New York 1944.)
- 1935b** Gestufte Räume; *Fundamenta Mathematicae* 25 (1935), S. 486 – 502.
- 1935c** Problem 62; *Fundamenta Mathematicae* 25 (1935), S. 578.
- 1936a** Über zwei Sätze von G. Fichtenholz und L. Kantorovitch; *Studia Mathematica* 6 (1936), S. 18 – 19.
- 1936b** Summen von $\aleph_1$ Mengen; *Fundamenta Mathematicae* 26 (1936), S. 241 – 255.
- 1937** Die schlichten stetigen Bilder des Nullraums; *Fundamenta Mathematicae* 29 (1937), S. 151 – 158.
- 1938** Erweiterung einer stetigen Abbildung; *Fundamenta Mathematicae* 30 (1938), S. 40 – 47.

3. Die mathematischen Schriften von Ernst Zermelo

- 1894** Untersuchungen zur Variationsrechnung; *Dissertation an der Universität Berlin*.
- 1896a** Über einen Satz der Dynamik und die mechanische Wärmetheorie; *Annalen der Physik und Chemie* 57 (1896), S. 485 – 494.
- 1896b** Über mechanische Erklärungen irreversibler Vorgänge. Eine Antwort auf Herrn Boltzmanns Entgegnung; *Annalen der Physik und Chemie* 59 (1896), S. 793 – 801.
- 1899** Über die Bewegung eines Punktsystems bei Bedingungsungleichungen; *Nachrichten von der Königlichen Gesellschaft der Wissenschaften zu Göttingen, mathematisch physikalische Klasse* (1899), S. 306 – 310.
- 1899/1900** Über die Anwendung der Wahrscheinlichkeitsrechnung auf dynamische Systeme; *Physikalische Zeitschrift* 1 (1899/1900), S. 317 – 320. (Habilitationsvortrag in Göttingen 1899.)
- 1901** Über die Addition transfiniter Kardinalzahlen. (Vorgelegt von D. Hilbert in der Sitzung vom 9. März 1901); *Nachrichten von der Königlichen Gesellschaft der Wissenschaften zu Göttingen, mathematisch physikalische Klasse* (1901), S. 34 – 38.
- 1902a** Hydrodynamische Untersuchungen über die Wirbelbewegungen in einer Kugelfläche. Erste Mitteilung; *Zeitschrift für Mathematik und Physik* 47 (1902), S. 201 – 237.
- 1902b** Zur Theorie der kürzesten Linien; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 11 (1902), S. 184 – 187.
- 1903** Über die Herleitung der Differentialgleichung bei Variationsproblemen; *Mathematische Annalen* 58 (1903), S. 558 – 564.
- 1904a** Beweis, daß jede Menge wohlgeordnet werden kann. (Aus einem an Herrn Hilbert gerichteten Briefe); *Mathematische Annalen* 59 (1904), S. 514 – 516.

- 1904b** Weiterentwicklung der Variationsrechnung in den letzten Jahren; *zusammen mit Hans Hahn. Enzyklopädie der mathematischen Wissenschaften mit Einschluß ihrer Anwendungen, Band 2, S. 626 – 641. Teubner, Leipzig.*
- 1908a** Neuer Beweis für die Möglichkeit einer Wohlordnung; *Mathematische Annalen* 65 (1908), S. 107 – 128.
- 1908b** Untersuchungen über die Grundlagen der Mengenlehre. I; *Mathematische Annalen* 65 (1908), S. 261 – 281.
- 1909a** Sur les ensembles finis et le principe de l'induction complète; *Acta Mathematica* 32 (1909), S. 186 – 193.
- 1909b** Über die Grundlagen der Arithmetik; *in: G. Castelnuovo (Hrsg.), Atti del IV Congresso Internazionale dei Matematici (Roma, 6 – 11 Aprile 1908), Band 2, S. 8 – 11, Accademia dei Lincei, Rom.*
- 1911** Die Einstellung der Grenzkonzentration an der Trennungsfläche zweier Lösungsmittel; *zusammen mit E. Riesenfeld. Physikalische Zeitschrift* 10 (1911), S. 958 – 961.
- 1913** Über eine Anwendung der Mengenlehre auf die Theorie des Schachspiels; *in: E. Hobson / A. Love (Hrsg), Proceedings of the 5th International Congress of Mathematics Cambridge 1912, Band 2, S. 501 – 504, Cambridge University Press, Cambridge.*
- 1914** Über ganze transzendente Zahlen; *Mathematische Annalen* 75 (1914), S. 434 – 442.
- 1927** Über das Maß und die Diskrepanz von Punktmengen; *Journal für die reine und angewandte Mathematik* 158 (1927), S. 154 – 167.
- 1928** Die Berechnung der Turnier-Ergebnisse als ein Maximumproblem der Wahrscheinlichkeitsrechnung; *Mathematische Zeitschrift* 29 (1928), S. 436 – 460.
- 1929** Über den Begriff der Definitheit in der Axiomatik; *Fundamenta Mathematicae* 14 (1929), S. 339 – 344.
- 1930a** Über Grenzzahlen und Mengenbereiche: Neue Untersuchungen über die Grundlagen der Mengenlehre; *Fundamenta Mathematicae* 16 (1930), S. 29 – 47.
- 1930b** Über die Navigation in der Luft als Problem der Variationsrechnung; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 39 (1930), S. 44 – 48.
- 1931** Über das Navigationsproblem bei ruhender oder veränderlicher Windverteilung; *Zeitschrift für angewandte Mathematik und Mechanik* 11 (1931), S. 114 – 124.
- 1932a** (Hrsg.) Georg Cantor – Gesammelte Abhandlungen mathematischen und philosophischen Inhalts; *Vorwort und Anmerkungen von Ernst Zermelo. Siehe Rubrik 5, „Gesammelte Werke“: Georg Cantor, 1932.*
- 1932b** Über die Stufen der Quantifikation und die Logik des Unendlichen; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 41 (1932), S. 85 – 88.
- 1932c** Über mathematische Systeme und die Logik des Unendlichen; *Forschungen und Fortschritte* 8 (1932), S. 6 – 7.
- 1933** Über die Bruchlinien zentrierter Ovale. Wie zerbricht ein Stück Zucker?; *Zeitschrift für angewandte Mathematik und Mechanik* 13 (1933), S. 168 – 170.
- 1934** Elementare Betrachtungen zur Theorie der Primzahlen; *Nachrichten von der Gesellschaft der Wissenschaften zu Göttingen, mathematisch physikalische Klasse* (1934), S. 43 – 46.
- 1935** Grundlagen einer allgemeinen Theorie der mathematischen Satzsysteme (erste Mitteilung); *Fundamenta Mathematicae* 25 (1935), S. 136 – 146.

4. Mathematische Schriften anderer Autoren

- Ackermann, Wilhelm** 1937 Die Widerspruchsfreiheit der allgemeinen Mengenlehre; *Mathematische Annalen* 114 (1937), S. 305 – 315.
- 1956 Zur Axiomatik der Mengenlehre; *Mathematische Annalen* 131 (1956), S. 336 – 345.
- Alexandrov, Pavel** 1916 Sur la puissance des ensembles mesurables B.; *Comptes Rendus Hebdomadaires des Séances de l'Académie des Sciences* 162 (1916), S. 323 – 325.
- Alexandrov, Pavel / Urysohn, Pavel** 1929 Mémoire sur les espaces topologiques compacts; *Verh. Akad. Wetensch., Amsterdam* 14 (1929), S. 1 – 96.
- Baire, René** 1899 Sur les fonctions de variables réelles; *Annali di Matematica Pura ed Applicata, Ser. IIIa* 3 (1899), S. 1 – 123.
- Banach, Stefan** 1923 Sur le problème de la mesure; *Fundamenta Mathematicae* 4 (1923), S. 7 – 33.
- 1924 Un théorème sur les transformations biunivoques; *Fundamenta Mathematicae* 6 (1924), S. 236 – 239.
- 1930 Über additive Maßfunktionen in abstrakten Mengen; *Fundamenta Mathematicae* 15 (1930), S. 97 – 101.
- Banach, Stefan / Kuratowski, Kazimierz** 1929 Sur une généralisation du problème de la mesure; *Fundamenta Mathematicae* 14 (1929), S. 127 – 131.
- Banach, Stefan / Tarski, Alfred** 1924 Sur la décomposition des ensembles de points en parties respectivement congruentes; *Fundamenta Mathematicae* 6 (1924), S. 244 – 277.
- Bendixson, Ivar** 1883 Quelques théorèmes de la théorie des ensembles de points; *Acta Mathematica* 2 (1883), S. 415 – 429.
- Bernays, Paul** 1937 – 1954 A system of axiomatic set theory. I – VII; jeweils in: *Journal of Symbolic Logic*. I: 2 (1937), S. 65 – 77, II: 6 (1941), S. 1 – 17, III: 7 (1942), S. 65 – 89, IV: 7 (1942), S. 133 – 145, V: 8 (1943), S. 89 – 106, VI: 13 (1948), S. 65 – 79, VII: 19 (1954), S. 81 – 96.
- Bernstein, Felix** 1905a Über die Reihe der transfiniten Ordnungszahlen; *Mathematische Annalen* 60 (1905), S. 187 – 193.
- 1905b Untersuchungen aus der Mengenlehre; *Mathematische Annalen* 61 (1905), S. 117 – 155. (Die Arbeit ist ein minimal veränderter Abdruck der Dissertation von Felix Bernstein, Göttingen 1901.)
- 1908 Zur Theorie der trigonometrischen Reihen; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 60 (1908), S. 325 – 338.
- Bloch, Gérard** 1953 Sur les ensembles stationnaires de nombres ordinaux et les suites distinguées de fonctions regressives; *Comptes Rendus Hebdomadaires des Séances de l'Académie des Sciences* 236 (1953), S. 265 – 268.
- Bois-Reymond, Paul du** 1882 Die allgemeine Funktionentheorie; *Tübingen*.
- Borel, Émile** 1898 Leçons sur la Théorie des Fonctions; *Gauthier-Villars, Paris*.
- Brouwer, Luitzen** 1911 Beweis der Invarianz der Dimensionszahl; *Mathematische Annalen* 70, S. 161 – 165.

- Burali-Forti, Cesare** 1897 Una questione sui numeri transfiniti; *Rendiconti del circolo matematico di Palermo* 11 (1897), S. 154 – 164.
- 1997b Sulle classi ben ordinate; *Rendiconti del circolo matematico di Palermo* 11 (1897), S. 260.
- Cohen, Paul** 1963 The independence of the continuum hypothesis. Part I; *Proceedings of the National Academy of Science USA* 50 (1963). S. 1143 – 1148.
- 1964 The independence of the continuum hypothesis. Part II; *Proceedings of the National Academy of Science USA* 51 (1964). S. 105 – 110.
- Dedekind, Richard** 1872 Stetigkeit und irrationale Zahlen; *Vieweg, Braunschweig*.
- 1888 Was sind und was sollen die Zahlen?; *Vieweg, Braunschweig* (8. Auflage 1960).
- Deiser, Oliver** 2007 Ordinalzahlen in der Analysis und Maßtheorie; *Mathematische Semesterberichte* 54/2 (2007), S. 177 – 197.
- Fodor, Geza** 1956 Eine Bemerkung zur Theorie der regressiven Funktionen; *Acta Scientiarum Mathematicarum, Szeged* 17 (1956), S. 139 – 142.
- Fraenkel, Abraham** 1922 Zu den Grundlagen der Cantor-Zermeloschen Mengenlehre; *Mathematische Annalen* 86 (1922), S. 230 – 237.
- 1927 Zehn Vorlesungen über die Grundlegung der Mengenlehre; *Leipzig*.
- Fréchet, Maurice** 1910 Les dimensions d'un ensemble abstrait; *Mathematische Annalen* 68 (1910), S. 145 – 168.
- Frege, Gottlob** 1879 Begriffsschrift. Eine der arithmetischen nachgebildete Formelsprache des reinen Denkens; *Louis Nebert, Halle*.
Daneben: Begriffsschrift und andere Aufsätze. Mit E. Husserls und H. Scholz' Anmerkungen herausgegeben von Ignacio Angelelli; *Georg Olms, Hildesheim* 1964.
- 1884 Die Grundlagen der Arithmetik. Eine logisch-mathematische Untersuchung über den Begriff der Zahl; *Wilhelm Koebner, Breslau*.
Neudruck bei M. & H. Marcus, Breslau 1934. *Reprographischer Nachdruck dieser Ausgabe bei Georg Olms, Hildesheim*, 1961. *Daneben eine von Christian Thiel herausgegebene und mit einer Einleitung versehene Ausgabe bei Felix Meiner, Hamburg*, 1988.
 - 1893 Grundgesetze der Arithmetik. Begriffsschriftlich abgeleitet von G. Frege. I; *Poble, Jena*. *Reprographischer Nachdruck bei Georg Olms, Hildesheim* 1962.
 - 1903 Grundgesetze der Arithmetik. Begriffsschriftlich abgeleitet von G. Frege. II; *Poble, Jena*. *Reprographischer Nachdruck bei Georg Olms, Hildesheim* 1962.
 - 1965 Funktion, Begriff, Bedeutung. Fünf logische Studien; *zweite Auflage*. Hrsg. Günther Patzig. *Vandenboeck und Ruprecht, Göttingen*.
 - 1966 Logische Untersuchungen; Hrsg. G. Patzig. *Vandenboeck und Ruprecht, Göttingen*.
 - 1976 Nachgelassene Schriften und wissenschaftlicher Briefwechsel; *Meiner, Hamburg*.
 - 1984 Collected Papers on mathematics, logic and philosophy; *Basil Blackwell, Oxford*.
- Friedman, Harvey** 1971 Higher set theory and mathematical practice; *Annals of Mathematical Logic* 2 (1971), S. 325 – 357.
- Gaifman, Haim** 1967 A generalization of Mahlo's method for obtaining large cardinal numbers; *Israel Journal of Mathematics* 5 (1967), S. 188 – 200.
- Gentzen, Gerhard** 1936 Die Widerspruchsfreiheit der reinen Zahlentheorie; *Mathematische Annalen* 112 (1936), S. 493 – 565.
- 1943 Beweisbarkeit und Unbeweisbarkeit von Anfangsfällen der transfiniten Induktion in der reinen Zahlentheorie; *Mathematische Annalen* 119 (1943), S. 140 – 161.

Gödel, Kurt 1930 Die Vollständigkeit der Axiome des logischen Funktionenkalküls; *Monatshefte für Mathematik und Physik* 37 (1930), S. 349 – 360.

- 1931 Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme; *Monatshefte für Mathematik und Physik* 38 (1931), S. 173 – 198.
- 1938 The consistency of the axiom of choice and the generalized continuum hypothesis; *Proceedings of the National Academy of Sciences USA* 24 (1938), S. 556 – 557.
- 1939 Consistency-proof for the Generalized Continuum-Hypothesis; *Proceedings of the National Academy of Sciences USA* 25 (1939), S. 220 – 224.
- 1940 The consistency of the Axiom of Choice and of the Generalized Continuum Hypothesis with the Axioms of Set Theory; *Annals of Mathematics Studies* 3, Princeton University Press, Princeton.

Greiling, Kurt / Nelson, Leonard 1908 Bemerkungen zu den Paradoxien von Russell und Burali-Forti. Mit Anhängen von H. Goesch und G. Hessenberg; *Abhandlungen der Friesschen Schule, Neue Folge, Band 2*, S. 301 – 331.

Halpern, Dan / Howard, Paul 1976 The law of infinite cardinal addition is weaker than the axiom of choice; *Transactions of the American Mathematical Society* 220 (1976), S. 195 – 204.

Hartogs, Friedrich 1915 Über das Problem der Wohlordnung; *Mathematische Annalen* 76, S. 438 – 443.

Hessenberg, Gerhard 1904 Das Unendliche in der Mathematik; *Abhandlungen der Friesschen Schule, Neue Folge, Band 1*, S. 137 – 190.

- 1906 Grundbegriffe der Mengenlehre; *Abhandlungen der Friesschen Schule, Neue Folge, Band 1*, S. 478 – 706. (In Buchform als Sonderdruck erschienen bei Vandenhoeck & Ruprecht, Göttingen, 1906.)
- 1907 Potenzen transfiniter Ordnungszahlen; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 16 (1907), S. 130 – 137.

Hermite, Charles 1874 Sur la fonction exponentielle; *Comptes Rendus Hebdomadaires des Séances de l'Académie des Sciences* 77 (1874), S. 18 – 24, 74 – 79, 226 – 233, 285 – 293.

Heyting, Arend 1930 Die formalen Regeln der intuitionistischen Mathematik; *Sitzungsberichte der Preussischen Akademie der Wissenschaften. Physikalisch-mathematische Klasse* (1930), S. 42 – 56 und S. 57 – 71.

Hilbert, David 1891 Über die stetige Abbildung einer Linie auf ein Flächenstück; *Mathematische Annalen* 38 (1891), S. 459 – 460.

- 1899 Grundlagen der Geometrie; in: *Festschrift zur Feier der Enthüllung des Gauss-Weber-Denkmal in Göttingen*, Teubner, Leipzig, S. 1 – 95. (2. erweiterte Auflage Leipzig 1903.)
- 1900 Mathematische Probleme. Vortrag, gehalten auf dem internationalen Mathematiker-Kongreß zu Paris 1900; *Nachrichten der Königl. Gesellschaft der Wissenschaften zu Göttingen, mathematisch-physikalische Klasse* (1900), S. 253 – 297. auch: *Archiv der Mathematik und Physik* 1 (1901), S. 44 – 63, 213 – 237.

Jacobsthal, Ernst 1908 Über den Aufbau der transfiniten Arithmetik; *Mathematische Annalen* 66 (1908), S. 145 – 194.

- 1909 Zur Arithmetik der transfiniten Zahlen; *Mathematische Annalen* 67 (1909), S. 130 – 144.

- Jensen, Ronald** 1968 Suslin's hypothesis is incompatible with $V = L$ (abstract); *Notices of the American Mathematical Society* 15 (1968), S. 935.
- 1972 The fine structure of the constructible hierarchy; *Annals of Mathematical Logic* 4 (1972), S. 229 – 308.
- Jech, Thomas** 1967 Nonprovability of Suslin's hypothesis; *Commentationes Mathematicae Universitatis Carolinae* 8 (1967), S. 291 – 305.
- Johansson, Ingebrigt** 1937 Der Minimalkalkül, ein reduzierter intuitionistischer Formalismus; *Compositio Mathematica* 4 (1937), S. 119 – 136.
- Jourdain, Philip** 1908 On the multiplication of alephs; *Mathematische Annalen* 65 (1908), S. 506 – 512.
- Kelley, John** 1975 General Topology; (*Nachdruck der Originalausgabe bei Van Nostrand, Toronto, Ont., 1955*). *Graduate Texts in Mathematics* 27, Springer, Berlin.
- König, Denes** 1927 Über eine Schlußweise aus dem Endlichen ins Unendliche: Punktmengen. Kartenfärben. Verwandtschaftsbeziehungen. Schachspiel; *Acta litterarum ac scientiarum regiae universitatis Hungaricae Francisco-Josephinae, sectio scientiarum mathematicarum* 3 (1927), S. 121 – 130.
- König, Julius** 1905a Zum Kontinuum-Problem. (Abgedruckt aus den Verhandlungen des III. Internationalen Mathematiker-Kongresses zu Heidelberg 1904.); *Mathematische Annalen* 60 (1905), S. 177 – 180.
- 1905b Über die Grundlagen der Mengenlehre und das Kontinuumproblem; *Mathematische Annalen* 61 (1905), S. 156 – 160.
 - 1906 Sur la théorie des ensembles; *Comptes Rendus Hebdomadaires des Séances de l'Académie des Sciences* 143, S. 110 – 112.
 - 1907 Über die Grundlagen der Mengenlehre und das Kontinuumproblem. Zweite Mitteilung; *Mathematische Annalen* 63 (1907), S. 217 – 221.
 - 1914 Neue Grundlagen der Logik, Arithmetik und Mengenlehre; *Veit, Leipzig*.
- Kolmogorov, Andrej** 1925 (On the principle of excluded middle); *Engl. Übersetzung des russischen Originals von 1925 in: Heijenoort (Hrsg.) 1967*.
- Korselt, Alwin** 1905 Über die Grundlagen der Mathematik; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 14 (1905), S. 365 – 389.
- 1906 Paradoxien der Mengenlehre; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 15 (1906), S. 215 – 219.
 - 1911 Über einen Beweis des Äquivalenzsatzes; *Mathematische Annalen* 70 (1911), S. 294 – 296.
- Kreisel, Georg / Levy, Azriel** 1968 Reflection principles and their use for establishing the complexity of axiomatic systems; *Zeitschrift für mathematische Logik* 14 (1968), S. 97 – 191.
- Kurata, Reijiro** 1964 On the existence of a proper complete model of set theory; *Commentarii Mathematici Universitatis Sancti Pauli* 13 (1964), S. 35 – 43.
- Kuratowski, Kazimierz** 1921 Sur la notion de l'ordre dans la théorie des ensembles; *Fundamenta Mathematicae* 2 (1921), S. 161 – 171.
- 1922 Une méthode d'élimination des nombres transfinis des raisonnements mathématique; *Fundamenta Mathematicae* 3 (1922), S. 76 – 108.
- Lebesgue, Henri** 1902 Intégrale, Longueur, Aire; *Thèse (1902), Paris*.

- 1905 Sur les fonctions représentables analytiquement; *Journal de Mathématiques Pures et Appliquées* 6 (1905), S. 139 – 216.
- 1907 Contribution à l'étude des correspondances de M. Zermelo; *Bulletin de la Société Mathématique de France* 35 (1907), S. 202 – 212.
- Levy, Azriel** 1969 The definability of cardinal numbers; in: *Foundation of Mathematics* (J. Bulloff, Hrsg.), Springer, Berlin. S. 15 – 38.
- Lew, John / Rosenberg, Arnold** 1978 Polynomial indexing of integer lattice points. I, II; *Journal of Number Theory* 10 (1978), S. 192 – 214, 215 – 243.
- Lindelöf, Ernst** 1905 Remarques sur un théorème fondamental de la théorie des ensembles; *Acta Mathematica* 29 (1905), S. 183 – 190.
- Lindemann, Ferdinand** 1882 Über die Zahl π ; *Mathematische Annalen* 20 (1882), S. 213 – 225.
- Lindenbaum, Adolf / Tarski, Alfred** 1926 Communication sur les recherches de la théorie des ensembles; *Comptes Rendus Hebdomadaires des Séances de la Société des Sciences et des Lettres de Varsovie XIX* (1926), classe III, S. 299 – 330.
- Liouville, Joseph** 1851 Sur des classes très-étendues de quantités dont la valeur n'est ni algébrique, ni même réductible à des irrationnelles algébriques; *Journal de mathématiques pures et appliquées* 16 (1851), S. 133 – 142.
- Lusin, Nikolai** 1917 Sur la classification de M. Baire; *Comptes Rendus Hebdomadaires des Séances de l'Académie des Sciences* 164 (1917), S. 91 – 94.
- 1925 a Sur un problème de M. Émile Borel et les ensembles projectifs de M. Henri Lebesgue; les ensembles analytiques; *Comptes Rendus Hebdomadaires des Séances de l'Académie des Sciences* 180 (1925), S. 1318 – 1320.
- 1925 b Sur les ensembles projectifs de M. Henri Lebesgue; *Comptes Rendus Hebdomadaires des Séances de l'Académie des Sciences* 180 (1925), S. 1572 – 1574.
- 1925 c Les propriétés des ensembles projectifs; *Comptes Rendus Hebdomadaires des Séances de l'Académie des Sciences* 180 (1925), S. 1817 – 1819.
- 1925 d Sur les ensembles non mesurables B et l'emploi de la diagonale Cantor; *Comptes Rendus Hebdomadaires des Séances de l'Académie des Sciences* 181 (1925), S. 95 – 96.
- 1927 Sur les ensembles analytiques; *Fundamenta Mathematicae* 10 (1927), S. 1 – 95.
- 1930 Leçons sur les ensembles analytiques et leurs applications; Gauthier-Villars, Paris. (Korrigierter Nachdruck: Chelsea, New York, 1972.)
- 1935 Sur les ensembles analytiques nuls; *Fundamenta Mathematicae* 25 (1935), S. 109 – 131.
- Lusin, Nikolai / Sierpiński, Waclaw** 1918 Sur quelques propriétés des ensembles (A); *Bulletin de l'Académie des Sciences Cracovie* (1918), S. 35 – 48.
- 1923 Sur un ensemble non mesurable B; *Journal de Mathématiques Pures et Appliquées* 9 (1923), S. 53 – 72.
- Mahlo, Paul** 1911 Über lineare transfinite Mengen; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 63 (1911), S. 187 – 225.
- 1912 Zur Theorie und Anwendung der $\aleph_0$ -Zahlen. I; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 64 (1912), S. 108 – 112.

- 1913 a Zur Theorie und Anwendung der $\aleph_0$ -Zahlen. II; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 65 (1913), S. 268 – 282.
- 1913 b Über Teilmengen des Kontinuums von dessen Mächtigkeit; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 65 (1913), S. 283 – 315.
- Martin, Donald** 1970 Measurable cardinals and analytic games; *Fundamenta Mathematicae* 66 (1970), S. 287 – 291.
- 1975 Borel determinacy; *Annals of Mathematics* 102(1975), S. 363 – 371.
- Martin, Donald / Steel, John** 1989 A proof of projective determinacy; *Journal of the American Mathematical Society* 2 (1989), S. 71 – 125.
- Méray, Charles** 1869 Remarques sur la nature des quantités définies par la condition de servir de limites à des variables données. *Revue des Sociétés savantes. Sciences mathém. phys. et naturelles, 2^e séries, IV* (1869).
- 1872 Nouveau précis d'analyse infinitésimale; *Paris*.
- Mirimanov, Dmitry** 1917a Les antinomies de Russell et de Burali-Forti et le problème fondamental de la théorie des ensembles; *L'Enseignement Mathématique* 19 (1917), S. 37 – 52.
- 1917b Remarques sur la théorie des ensembles et les antinomies cantoriniennes. I; *L'Enseignement Mathématique* 19 (1917), S. 209 – 217.
- 1919 Remarques sur la théorie des ensembles et les antinomies cantoriniennes. II; *L'Enseignement Mathématique* 21 (1919), S. 29 – 52.
- Mostowski, Andrzej** 1951 Some impredicative definitions in axiomatic set theory; *Fundamenta Mathematicae* 37 (1951), S. 111 – 124, Korrektur: 38 (1952), S. 238.
- Mycielski, Jan / Steinhaus, Hugo** 1962 A mathematical axiom contradicting the axiom of choice; *Bulletin de l'Académie Polonaise des Sciences* 10 (1962), S. 1 – 3.
- Neumann, John von** 1923 Zur Einführung der transfiniten Zahlen; *Acta Litterarum ac Scientiarum Regiae Universitatis Hungaricae Francisco-Josephinae, Sectio Scientiarum Mathematicarum* 1922/1923, Szeged, S. 199 – 208.
- 1925 Eine Axiomatisierung der Mengenlehre; *Journal für die reine und angewandte Mathematik* 154 (1925), S. 219 – 240.
- 1928a Die Axiomatisierung der Mengenlehre; *Mathematische Zeitschrift* 27, (1928), S. 669 – 752.
- 1928b Über die Definition durch transfinite Induktion und verwandte Fragen der allgemeinen Mengenlehre; *Mathematische Annalen* 99, (1928), S. 373 – 391.
(Mit einem sich anschließenden Zusatz von Abraham Fraenkel.)
- Peano, Giuseppe** 1888 Calcolo geometrico secondo l'Ausdehnungslehre di H. Grassmann, proceduto dalle operazioni della logica deduttiva; *Turin*.
- 1889 Arithmetices principia, nova methodo exposita; *Turin*.
- 1890 Sur une courbe qui remplit une aire plane; *Mathematische Annalen* 36 (1890), S. 157 – 160.
- 1906 Super theoremata de Cantor-Bernstein; *Rendiconti del circolo matematico di Palermo* 21 (1906), S. 360 – 366. Daneben auch abgedruckt und mit einer „Addition“ versehen in: *Revista de matematica* 8 (1906), S. 136 – 143.

- Russell, Bertrand** 1902 Brief an Gottlob Frege vom 16.6.1902;
In: *Heijenoort (Hrsg.)* 1967, S. 124 – 125.
- 1903 *The Principles of Mathematics*; *Cambridge University Press, Cambridge*.
(Zweite Auflage 1937. Nachdruck: *George Allen & Unwin, London*, 1948.)
 - 1908 Mathematical logic as based on the theory of types; *American Journal of Mathematics* 30 (1908), S. 222 – 262.
- Russell, Bertrand / Whitehead, Alfred North** 1910, 1912, 1913 *Principia Mathematica* I – III; *Cambridge University Press, Cambridge*.
- Siegel, Carl** 1949 *Transcendental Numbers*; *Princeton University Press, Princeton*.
- Sageev, Gershon** 1975 An independence result concerning the axiom of choice;
Annals of Mathematical Logic 8 (1975), S. 1 – 184.
- Scheffler, Ludwig** 1884 Zur Theorie der stetigen Funktionen einer reellen Veränderlichen; *Acta Mathematica* 5 (1884), S. 279 – 296.
- Schröder, Ernst** 1873 Lehrbuch der Arithmetik und Algebra; (Band 1). *Leipzig*.
 - 1877 Der Operationskreis des Logikkalküls; *Leipzig*. Nachdruck: *Teubner, Stuttgart* 1966.
 - 1880 – 1905 Vorlesungen über die Algebra der Logik (Exakte Logik);
Drei Bände in fünf erschienenen Teilen. *Teubner, Leipzig*. Band I 1890, Band II.1 1891,
Band II.2 1905, Band III.1 1895. Nachdruck bei *Chelsea, New York* 1966 (zweite Auflage),
einschließlich des von Eugen Müller bearbeiteten zweiteilig erschienenen „Abriß der
Algebra der Logik“, *Teubner, Leipzig* 1909 (Teil 1), 1910 (Teil 2).
- Shepherdson, John** 1951, 1952, 1953 Inner models for set theory; *Journal of Symbolic Logic* 16 (1951), S. 161 – 190; 17 (1952), S. 225 – 237; 18 (1953), S. 145 – 167.
- Shoenfield, Joseph** 1954 A relative consistency proof; *Journal of Symbolic Logic* 19 (1954), S. 21 – 28.
- Sierpiński, Waclaw** 1925 Sur une classe d'ensembles; *Fundamenta Mathematicae* 7 (1925), S. 237 – 243.
 - 1927 Sur quelques propriétés des ensembles projectifs; *Comptes Rendus Hebdomadaires des Séances de l'Académie des Sciences* 185 (1927), S. 833 – 835.
 - 1928 Sur une décomposition d'ensembles; *Monatshefte für Mathematik und Physik* 35 (1928), S. 239 – 242.
 - 1930 Sur l'uniformisation des ensembles mesurables; *Fundamenta Mathematicae* 16 (1930), S. 136 – 139.
- Sierpiński, Waclaw / Tarski, Alfred** 1930 Sur une propriété caractéristique des nombres inaccessibles; *Fundamenta Mathematicae* 15 (1930), S. 292 – 300.
- Skolem, Thoralf** 1923 Einige Bemerkungen zur axiomatischen Begründung der Mengenlehre; *Wissenschaftliche Vorträge gehalten auf dem Fünften Kongreß der Skandinavischen Mathematiker in Helsingfors vom 4. – 7. Juli 1922*. *Akademische Buchhandlung, Helsingfors* 1923. S. 217 – 232.
 - 1929 Über die Grundlagendiskussion in der Mathematik; *Proceedings of the 7th Scandinavian Mathematical Congress, Oslo* 1929, S. 3 – 21.
- Solovay, Robert** 1971 Real-valued measurable cardinals; in: *Dana Scott (Hrsg.): Axiomatic Set Theory. Proceedings of Symposia in Pure Mathematics. Vol. 13*. *American Mathematical Society, Providence*, 1971, S. 397 – 428.
- Solovay, Robert / Tennenbaum, Stanley** 1971 Iterated Cohen extensions and Suslin's problem; *Annals of Mathematics* 94 (1971), S. 201 – 245.

- Suslin, Mikhail** 1917 Sur une définition des ensembles mesurables B sans nombres transfinis; *Comptes Rendus Hebdomadaires des Séances de l'Académie des Sciences* 164 (1917), S. 88 – 91.
- 1920 Problème 3; *Fundamenta Mathematicae* 1 (1920), S. 223.
- Tarski, Alfred** 1924a Sur quelques théorèmes qui équivalent à l'axiome du choix; *Fundamenta Mathematicae* 5 (1924), S. 147 – 154.
- 1924b Sur les ensembles finis; *Fundamenta Mathematicae* 6 (1924), S. 45 – 95.
- 1925 Quelques théorèmes sur les alephs; *Fundamenta Mathematicae* 7 (1925), S. 1 – 14.
- 1930 Une contribution à la théorie de mesure; *Fundamenta Mathematicae* 15 (1930), S. 42 – 50.
- 1935 Der Wahrheitsbegriff in den formalisierten Sprachen; *Studia Philosophica* 1, S. 261 – 405 (1935). (Deutsche Übersetzung der polnischen Arbeit von 1933.)
- 1962 Some problems and results relevant to the foundations of set theory; in: *Logic, Methodology and Philosophy of Science* (Hrsg. Ernst Nagel, Patrick Suppes, Alfred Tarski), S. 125 – 135, Stanford University Press, Stanford.
- Teichmüller, Oswald** 1939 Braucht der Algebraiker das Auswahlaxiom?; *Deutsche Mathematik* 4 (1939), S. 567 – 577.
- Tennenbaum, Stanley** 1968 Suslin's Problem; *Proceedings of the National Academy of the U.S.A.* 59 (1968), S. 60 – 63.
- Tukey, John** 1940 Convergence and Uniformity in Topology; Princeton University Press, Princeton.
- Ulam, Stanislaw** 1930 Zur Maßtheorie der allgemeinen Mengenlehre; *Fundamenta Mathematicae* 16 (1930), S. 140 – 150.
- Veblen, Oswald** 1908 Continuous increasing functions of finite and transfinite ordinals; *Transactions of the American Mathematical Society* 9 (1908), S. 280 – 292.
- Venn, John** 1880 On the diagrammatic and mechanical representation of propositions and reasonings; *The London, Edinburgh and Dublin Philosophical Magazine and Journal of Science* 10 (1880), S. 1 – 18.
- Vitali, Giuseppe** 1905 Sul problema della misura dei gruppi di punti di una retta; *Gamberini e Parmeggiani, Bologna*.
- Waerden, Bartel Leendert van der** 1930 Moderne Algebra; Springer, Berlin.
- Wang, Hao** 1949 On Zermelo's and von Neumann's axioms for set theory; *Proceedings of the National Academy of Sciences* 35 (1949), S. 150 – 155.
- Weber, Heinrich** 1906 Elementare Mengenlehre; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 15 (1906), S. 173 – 184.
- Woodin, W. Hugh** 2001 The Continuum Hypothesis, Part I and II; *Notices of the American Mathematical Society* 48 (2001), S. 567 – 576 (Teil I), 681 – 690 (Teil 2).
- Young, William** 1903 Zur Lehre der nicht abgeschlossenen Punktmengen; *Berichte über die Verhandlungen der Königlich Sächsischen Gesellschaft der Wissenschaften zu Leipzig, Mathematisch-Physikalische Klasse* 55 (1903), S. 287 – 293.
- Zorn, Max** 1935 A remark on method in transfinite algebra; *Bulletin of the American Mathematical Society* 41 (1935), S. 667 – 670.
- 1944 Idempotency of infinite cardinals; *University of California Publications* 2, S. 9 – 12.

5. Gesammelte Werke, Briefe und Aufsatzsammlungen

- Borel, Émile** 1972 *Œuvres de Émile Borel; vier Bände. Editions du Centre National de la Recherche Scientifique, Paris.*
- Cantor, Georg** 1932 *Gesammelte Abhandlungen mathematischen und philosophischen Inhalts. Mit erläuternden Anmerkungen sowie mit Ergänzungen aus dem Briefwechsel Cantor-Dedekind. Herausgegeben von Ernst Zermelo. Springer, Berlin. Reprographischer Nachdruck Georg Olms Verlagsbuchhandlung, Hildesheim 1962. Bibliographisch ergänzte Neuauflagen 1980, 1990 bei Springer, Berlin.*
- 1984 *Über unendliche, lineare Punktmannigfaltigkeiten. Arbeiten zur Mengenlehre aus den Jahren 1872 – 1884; Herausgegeben von Günter Asser, B. G. Teubner, Leipzig.*
 - 1991 *Briefe; Hrsg. Herbert Meschkowski und Winfried Nilson. Springer, Berlin.*
- Cantor, Georg / Dedekind, Richard** 1937 *Briefwechsel Cantor – Dedekind; Herausgegeben von E. Noether und J. Cavaillès. Hermann, Paris.*
- Dedekind, Richard** 1930 – 1932 *Gesammelte mathematische Werke; drei Bände. Herausgegeben von Robert Fricke, Emmy Noether und Øystein Ore. Vieweg, Braunschweig. (Nachdruck in zwei Bänden: Chelsea, New York, 1969).*
- Ewald, William** (Hrsg.) 1996 *From Kant to Hilbert. A Source Book in the Foundations of Mathematics; in zwei Bänden. Clarendon Press, Oxford.*
- Felgner, Ulrich** (Hrsg.) 1979 *Mengenlehre; Wissenschaftliche Buchgesellschaft, Darmstadt.*
- Finsler, Paul** 1979 *Aufsätze zur Mengenlehre; Herausgegeben von Georg Unger. Wissenschaftliche Buchgesellschaft, Darmstadt.*
- Gödel, Kurt** 1986, 1990, 1995, 2003, 2003b *Collected Works, Volume I – V; Oxford University Press, Oxford.*
- Hilbert, David** 1935 *Gesammelte Abhandlungen; Band 3: Analysis, Grundlagen der Mathematik, Physik, Verschiedenes, Lebensgeschichte. Springer, Berlin. 2. Auflage 1970.*
- Hausdorff, Felix** 1969 *Nachgelassene Schriften in zwei Bänden. Studien und Referate; (enthält Handschriften-Faksimiles). Herausgegeben von Günter Bergmann. Teubner, Stuttgart.*
- Hausdorff, Felix** 2001 ff. *Gesammelte Werke in 9 Bänden; Springer, Berlin.*
- Heijenoort, Jean van** (Hrsg.) 1967 *From Frege to Gödel. A Source Book in Mathematical Logic, 1879 – 1931; Harvard University Press, Cambridge, Mass.*
- Kuratowski, Kazimierz** 1988 *Selected Papers; Hrsg. Karol Borsuk, Ryszard Engelking, Czesław Ryll-Nerdzewski. PWN – Polish Scientific Publishers, Warsaw.*
- Lebesgue, Henri** 1972 – 1973 *Œuvres Scientifiques. En Cinq Volumes; fünf Bände. L'Enseignement Mathématique, Genf.*
- Mostowski, Andrzej** 1979 *Foundational Studies. Selected Works; zwei Bände. Hrsg. Kazimierz Kuratowski u.a. PWN – Polish Scientific Publishers, Warsaw und North-Holland, Amsterdam.*
- Neumann, John von** 1961 *Collected Works, Volume I: Logic, Theory of Sets, and Quantum Mechanics; Hrsg. A. H. Taub. Pergamon Press, Oxford.*
- Peano, Giuseppe** 1957 – 1959 *Opere Scelte; drei Bände. Cremonese, Rom.*

- Skolem, Thoralf** 1970 Selected Works in Logic; Hrsg. Jens Erik Fenstad. *The Norwegian Research Council for Science and the Humanities, Scandinavian University Books, Universitetsforlaget, Oslo.*
- Tarski, Alfred** 1986 Collected Papers; 4 Bände. Hrsg. S.R. Givant und R.N. Mc Kenzie. *Birkhäuser, Basel.*

6. Ältere Bücher zur Mengenlehre

- Baire, René** 1905 Leçons sur les fonctions discontinues; *Gauthier-Villars, Paris.*
- Borel, Émile** 1898 Leçons sur la théorie des fonctions; *Gauthier-Villars, Paris.*
- Fraenkel, Abraham** 1919 Einleitung in die Mengenlehre; Eine allgemeinverständliche Einführung in das Reich der unendlichen Größen; *Springer, Berlin* 1919
- 1923 Einleitung in die Mengenlehre; 2. erweiterte Auflage, *Springer, Berlin.*
 - 1928 Einleitung in die Mengenlehre; 3. umgearbeitete und stark erweiterte Auflage, *Springer, Berlin.*
 - 1959 Mengenlehre und Logik; *Dunker & Humblot, Berlin.*
- Grelling, Kurt** 1924 Mengenlehre; *Mathematisch-Physikalische Bibliothek* 58, *Teubner, Leipzig.*
- Hausdorff, Felix** 1914 Grundzüge der Mengenlehre; *Veit & Comp., Leipzig.* (Nachdrucke bei *Chelsea, New York* 1949, 1965, 1978.)
- 1927 Mengenlehre. Zweite, neubearbeitete Auflage; zweite, stark umgearbeitete Auflage von „Grundzüge der Mengenlehre“, 1914. *Götschen Lehrbücher. 1. Gruppe: Reine Mathematik. Band 7. Walter de Gruyter, Berlin.*
 - 1935 Mengenlehre; dritte Auflage. *Walter de Gruyter, Berlin.* (Nachdruck bei *Dover, New York* 1944).
 - 1957 Set Theory; *Chelsea Publishing Company, New York.* (Englische Übersetzung von „Mengenlehre“, 1937.)
 - 2002 Gesammelte Werke, Band II: Grundzüge der Mengenlehre; *Springer, Berlin.* Die Ausgabe enthält eine Faksimile-Wiedergabe der „Grundzüge“ von 1914, eine historische Einführung und kommentierende Essays.
- Hessenberg, Gerhard** 1906 Grundbegriffe der Mengenlehre; *Vandenboeck & Ruprecht, Göttingen.* (Sonderdruck aus den *Abhandlungen der Friesschen Schule, Neue Folge*, 1. Band, 4. Heft, S. 478 – 706.)
- Hobson, Ernest** 1907 The Theory of Functions of a Real Variable and the Theory of Fourier's Series; *Cambridge University Press, Cambridge.*
- Kamke, Erich** 1928 Mengenlehre; *Sammlung Götschen* 999, *Walter de Gruyter, Berlin.* Viele weitere Auflagen.
- Lusin, Nikolai** 1930 Leçons sur les Ensembles Analytiques et Leurs Applications. Avec une Note de M. Sierpiński. Préface de M. Henri Lebesgue; *Gauthier-Villars, Paris.*
- Schoenflies, Arthur** 1900 Die Entwicklung der Lehre von den Punktmannigfaltigkeiten. Bericht, erstattet der Deutschen Mathematiker-Vereinigung; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 8, Heft 2, S. I – VI, 1 – 251, *B. G. Teubner, Leipzig.*

- 1908 Die Entwicklung der Lehre von den Punktmannigfaltigkeiten. Teil II; *Jahresbericht der Deutschen Mathematiker-Vereinigung*, 2. Ergänzungsband; B. G. Teubner, Leipzig.
- 1913 Entwicklung der Mengenlehre und ihrer Anwendungen. Erste Hälfte: Allgemeine Theorie der unendlichen Mengen und Theorie der Punktmengen. Umarbeitung des im VIII. Bande der Jahresberichte der Deutschen Mathematiker-Vereinigung erstatteten Berichts; B. G. Teubner, Leipzig.
- Shegalkin, Iwan** 1907 Transfinite Zahlen (russisch); *Universitätsdruckerei, Moskau*.
- Sierpiński, Waclaw** 1912 Zarys teorii mnogości; *Biblioteka matematyczno-fizyczna, Serie III, Band 9, Warschau*.
- 1928 Leçons sur les nombres transfinis; *Gauthier-Villars, Paris*.
- 1934 Hypothèse du continua; *Monografie Matematyczne IV, Warschau*.
- Young, William / Chisholm-Young, Grace** 1906 The Theory of Sets of Points; *Cambridge University Press, Cambridge*.

7. Bücher zur Mengenlehre bis etwa 1980

- Abian, Alexander** 1965 The Theory of Sets and Transfinite Arithmetic; *W.B. Saunders, Philadelphia, 1965*.
- Alexandrov, Pavel** 1984 Einführung in die Mengenlehre und in die allgemeine Topologie; *Übersetzung aus dem Russischen von Manfred Peschel u.a.. VEB Deutscher Verlag der Wissenschaften, Berlin*.
- Asser, Günter** 1988 Grundbegriffe der Mathematik. I. Mengen. Abbildungen. Natürliche Zahlen; *Fünfte Auflage. VEB Deutscher Verlag der Wissenschaften, Berlin*.
- Bachmann, Heinz** 1967 Transfinite Zahlen; *zweite neubearbeitete Auflage, Springer, Berlin. Erste Auflage 1955*.
- Bernays, Paul** 1958 Axiomatic Set Theory. With a historical introduction by Abraham Fraenkel; *North-Holland, Amsterdam. Zweite Auflage 1968. Reprint der zweiten Auflage: Dover, New York, 1991*.
- Bourbaki, Nicolas** 1970 Théorie des ensembles; *Éléments de mathématique, Hermann, Paris. (Ursprünglich begonnen 1939, danach verschiedene Bände und Neuauflagen.)*
- Breuer, Josef** 1972 Einführung in die Mengenlehre; *sechste Auflage, Schroedel, Hannover; Schöningh Paderborn*.
- Cohen, Paul** 1966 Set Theory and the Continuum Hypothesis; *W.A. Benjamin, New York*.
- Enderton, Herbert** 1977 Elements of Set Theory; *Academic Press, New York*.
- Halmos, Paul Richard** 1960 Naive Set Theory; *Van Nostrand, Princeton, NJ*.
- 1976 Naive Mengenlehre; *Vierte Auflage. Aus dem Amerikanischen übersetzt von Manfred Armbrust und Fritz Ostermann. (Deutsche Übersetzung von „Naive Set Theory“, 1960.) Vandenhoeck & Ruprecht, Göttingen*.
- Hasse, Maria** 1967 Grundbegriffe der Mengenlehre und Logik; *dritte erweiterte Auflage. Mathematische Schülerbücherei 2, Teubner, Leipzig*.
- Hayden, Seymour / Kennison, John** 1968 Zermelo-Fraenkel Set Theory; *Charles E. Merrill, Columbus, Ohio*.

- Jensen, Ronald** 1967 Modelle der Mengenlehre. Widerspruchsfreiheit und Unabhängigkeit der Kontinuumshypothese und des Auswahlaxioms. Ausgearbeitet von Franz Josef Leven; *Lecture Notes in Mathematics* 37, Springer, Berlin.
- Klaue, Dieter** 1979 Mengenlehre; *Walter de Gruyter, Berlin*.
- Krivine, Jean-Louis** 1969 Théorie axiomatique des ensembles; *Presses Universitaires de France, Paris*.
- Kuratowski, Kazimierz** 1972 Introduction to Set Theory and Topology; *Zweite Auflage; Pergamon Press, Oxford*.
- Kuratowski, Kazimierz / Mostowski, Andrzej** 1976 Set Theory; *zweite überarbeitete Auflage (erste Auflage 1968, engl. Übersetzung der polnischen Ausgabe von 1966). North-Holland, Amsterdam*.
- Monk, Donald** 1969 Introduction to Set Theory; *McGraw-Hill, New York*.
- Morse, Anthony** 1965 A Theory of Sets; *Academic Press, New York. Zweite Auflage (mit einem Vorwort von Trevor McNinn): Academic Press, Orlando, Florida, 1986*.
- Pinter, Charles** 1971 Set Theory; *Addison-Wesley, Reading, Mass.*
- Rotman, Brian / Kneebone, G. T.** 1966 The Theory of Sets and Transfinite Numbers; *Oldbourne, London*.
- Rubin, Jean** 1966 Set Theory for the Mathematician; *Holden-Day, San Francisco*.
- Schmidt, Jürgen** 1974 Mengenlehre. Band 1: Grundbegriffe; 2. verbesserte und erweiterte Auflage. *Bibliographisches Institut, Mannheim*.
- Sierpiński, Waclaw** 1965 Cardinal and Ordinal Numbers; *Warschau*.
- Skolem, Thoralf** 1962 Abstract Set Theory; *Notre Dame Mathematical Lectures* 8, *University of Notre Dame Press, Notre Dame*.
- Suppes, Patrick** 1960 Axiomatic Set Theory; *Van Nostrand, Princeton*. (Erweiterte und korrigierte Auflage bei Dover, New York, 1972.)
- Takeuti, Gaisi / Zaring, Wilson** 1971 Introduction to Axiomatic Set Theory; *Graduate Texts in Mathematics* 1, Springer, Berlin.
- Quine, Willard van Orman** 1969 Set Theory and its Logic; *Belknap Press of Harvard University Press, Cambridge, Mass. Zweite Auflage. (Erste Auflage 1963.) Deutsche Ausgabe: „Mengenlehre und ihre Logik“, Übersetzung von Anneliese Oberschelp. Logik und Grundlagen der Mathematik 10. Vieweg, Braunschweig 1973*.

8. Neuere Bücher zur Mengenlehre

- Deiser, Oliver** 2007 Reelle Zahlen. Das klassische Kontinuum und die natürlichen Folgen; *Springer, Berlin*.
- Devlin, Keith** 1993 The Joy of Sets – Fundamentals of Contemporary Set Theory; 2. Auflage. *Undergraduate Texts in Mathematics, Springer, New York*.
- Drake, Frank / Singh, Dasharath** 1996 Intermediate Set Theory; *John Wiley, Chichester*.
- Ebbinghaus, Heinz-Dieter** 2003 Einführung in die Mengenlehre; 4. Auflage. *Spektrum Akademischer Verlag, Heidelberg. (Frühere Auflagen bei Wissenschaftliche Buchgesellschaft, Darmstadt.)*

- Friedrichsdorf, Ulf / Prestel, Alexander** 1985 Mengenlehre für den Mathematiker; *Vieweg Studium: Grundkurs Mathematik*, Vieweg, Braunschweig.
- Hajnal, András / Hamburger, Peter** 1999 Set Theory; *Cambridge University Press, Cambridge*. (Übersetzung der ungarischen Originalausgabe von 1983 durch Attila Máté.)
- Hrbacek, Karel / Jech, Thomas** 1999 Introduction to Set Theory; 3. Auflage, *Monographs and Textbooks in Pure and Applied Mathematics 220*, Marcel Dekker, New York.
- Jech, Thomas** 1978 Set Theory; *Academic Press, New York*.
– 2003 Set Theory. The Third Millennium Edition, Revised and Expanded; *Springer, Berlin*.
- Just, Winfried / Weese, Martin** 1996, 1997 Discovering Modern Set Theory I & II: Set-theoretic Tools for Every Mathematician; *American Mathematical Society, Providence*.
- Kunen, Kenneth** 1980 Set Theory – An Introduction to Independence Proofs; *Studies in Logic and the Foundations of Mathematics Vol. 102*, North-Holland, Amsterdam.
- Levy, Azriel** 1979 Basic Set Theory; *Perspectives in Mathematical Logic*, Springer, Berlin.
- Moschovakis, Yiannis** 1994 Notes on Set Theory; *Springer, New York*.
- Oberschelp, Arnold** 1994 Allgemeine Mengenlehre; *Bibliographisches Institut, Mannheim*.
- Potter, Michael** 1990 Sets. An Introduction; *Clarendon Press, Oxford*. (Deutsche Übersetzung „Mengenlehre“, übersetzt von Achim Wittmüß, Spektrum Akademischer Verlag, Heidelberg, 1994.)
- Rautenberg, Wolfgang** 2007 Messen und Zählen. Eine einfache Konstruktion der reellen Zahlen; *Heldermann, Berlin*.
- Tourlakis, George** 2003 Lectures in Logic and Set Theory; 2 Bände, *Cambridge University Press, Cambridge*.
- Vaught, Robert** 1995 Set Theory; *zweite Auflage (erste Aufl. 1985)*. Birkhäuser, Boston.

9. Bücher zu speziellen Themen der Mengenlehre

- Aczel, Peter** 1988 Non-Well-Founded Sets; *CSLI Lecture Notes 14*, Stanford.
- Bartoszynsky, Tomek / Judah, Haim** 1995 Set Theory – On the Structure of the Real Line; *Peters, Wellesley, Mass.*
- Barwise, John** 1975 Admissible Sets and Structures; *Perspectives in Mathematical Logic*. Springer, Berlin.
- Bell, John** 1985 Boolean-valued models and independence proofs in set theory. With a foreword by Dana Scott; 2. Auflage. *Clarendon Press, Oxford*. (Erste Auflage 1977.)
- Dales, H. Garth / Woodin, W. Hugh** 1987 Independence for Analysts; *London Mathematical Society Lecture Note Series 115*, Cambridge University Press, Cambridge.
- Dehornoy, Patrick** 2000 Braids and Self-Distributivity; *Birkhäuser, Basel*.
- Devlin, Keith** 1984 Constructibility; *Perspectives in Mathematical Logic*. Springer, Berlin.
- Dodd, Anthony** 1982 The Core Model; *London Mathematical Society Lecture Note Series 61*, Cambridge University Press, Cambridge.
- Drake, Frank** 1974 Set Theory – An Introduction to Large Cardinals; *North-Holland, Amsterdam*.

- Erdős, Paul / Hajnal, András / Máté, Attila / Rado, Richard** 1984 Combinatorial Set Theory: Partition Relations for Cardinals; *North-Holland, Amsterdam*.
- Felgner, Ulrich** (Hrsg.) 1971 Models of ZF-Set Theory; *Springer Lecture Notes in Mathematics* 223, Springer, Berlin.
- Foreman, Matthew / Kanamori, Akihiro / Magidor, Menachem** (Hrsg.) 200? Handbook of Set Theory.
- Foster, Thomas** 1995 Set Theory with a Universal Set; *Clarendon Press, Oxford*.
- Friedman, Sy** 2000 Fine Structure and Class Forcing; *Walter de Gruyter, Berlin*.
- Herrlich, Horst** 2006 Axiom of Choice; *Lecture Notes in Mathematics*. Springer, Berlin.
- Howard, Paul / Rubin, Jean** 1998 Consequences of the Axiom of Choice; *American Mathematical Society, Providence*.
- Jech, Thomas** 1973 The Axiom of Choice; *North-Holland, Amsterdam*.
– 1986 Multiple Forcing; *Cambridge University Press, Cambridge*.
- Kanamori, Akihiro** 1994 The Higher Infinite. Large cardinals in set theory from their beginnings; *Perspectives in Mathematical Logic*, Springer, Berlin.
- Kechris, Alexander** 1995 Classical Descriptive Set Theory; *Graduate Texts in Mathematics* 156, Springer, New York.
- Mansfield, Richard / Weitkamp, Galen** 1985 Recursive Aspects of Descriptive Set Theory; *Clarendon Press, Oxford*.
- Martin, Donald** 200? Buch über das Axiom der Determiniertheit (englisch).
- Mitchell, William / Steel, John** 1994 Fine Structure and Iteration Trees; *Springer, Berlin*.
- Moschovakis, Yiannis** 1980 Descriptive Set Theory; *North-Holland, Amsterdam*.
- Neeman, Itay** 2004 The Determinacy of Long Games ; *Walter de Gruyter, Berlin*.
- Rubin, Herman / Rubin, Jean** 1963, 1985 Equivalents of the Axiom of Choice I & II; *North-Holland, Amsterdam*.
- Shelah, Saharon** 1994 Cardinal Arithmetic; *Clarendon Press, Oxford*.
– 1998 Proper and Improper Forcing; *Perspectives in Mathematical Logic*, Springer, Berlin.
- Steel, John** 1996 The Core Model Iterability Problem; *Springer, Berlin*.
- Wagon, Stanley** 1993 The Banach-Tarski Paradox; *zweite Auflage*. Cambridge University Press, Cambridge.
- Woodin, W. Hugh** 1999 The Axiom of Determinacy, Forcing Axioms and the Nonstationary Ideal; *Walter de Gruyter, Berlin*.
- Zeman, Martin** 2001 Inner Models and Large Cardinals; *Walter de Gruyter, Berlin*.

10. Bücher zur mathematischen Logik

- Asser, Günther** 1972, 1975, 1981 Einführung in die mathematische Logik (in drei Teilen, Neuauflagen); *Teubner, Leipzig*.
- Barwise, Jon** (Hrsg.) 1977 Handbook of Mathematical Logic; *North-Holland, Amsterdam*

- Boolos, George / Burgess, John / Jeffrey, Richard** 2002 Computability and Logic; 4. Auflage. Cambridge University Press, Cambridge.
- Chang, Chen-Chung / Keisler, H. Jerome** 1990 Model Theory; 3. aktualisierte Auflage (erste Auflage 1973), North-Holland, Amsterdam.
- Church, Alonzo** 1956 Introduction to Mathematical Logic; Princeton Univ. Press.
- Cutland, Nigel** 1980 Computability; Cambridge University Press, Cambridge.
- Ebbinghaus, Heinz-Dieter / Flum, Jörg / Thomas, Wolfgang** 1996 Einführung in die mathematische Logik; Vierte Auflage. Spektrum Akademischer Verlag, Heidelberg. (Englische Übersetzung "Mathematical Logic", 2. Auflage, Springer 1994.)
- Ebbinghaus, Heinz-Dieter / Flum, Jörg** 1999 Finite Model Theory; Zweite Auflage, Springer, Berlin.
- Gödel, Kurt** 1983 The Incompleteness Phenomenon; Peters, Wellesley, Mass.
- Hájek, Petr / Pudlák, Pavel** 1993 Metamathematics of First-Order Arithmetic; Perspectives in Mathematical Logic, Springer, Berlin.
- Hermes, Hans** 1976 Einführung in die mathematische Logik. Klassische Prädikatenlogik; 4. Auflage, Teubner, Stuttgart.
- Hinman, Peter** 2005 Fundamentals of Mathematical Logic; Peters, Wellesley, Mass.
- Hodges, Wilfrid** 1997 A Shorter Model Theory; Cambridge University Press, Cambridge.
- Kleene, Stephen Cole** 1967 Mathematical Logic; John Wiley, New York.
- Malitz, Jerome** 1987 Introduction to Mathematical Logic; 2. Auflage. Springer, Berlin.
- Manin, Yuri** 1977 A Course in Mathematical Logic; Englische Übersetzung aus dem Russischen von Neal Koblitz. Graduate Texts in Mathematics 53, Springer, Berlin.
- Marker, David** 2002 Model Theory. An Introduction; Springer, Berlin.
- Mates, Benson** 1997 Elementare Logik. Prädikatenlogik der ersten Stufe; Vandenhoeck & Ruprecht, Göttingen.
- Monk, Donald** 1976 Mathematical Logic, Model Theory. Springer, Berlin.
- Prestel, Alexander** 1986 Einführung in die mathematische Logik und Modelltheorie; Vieweg, Braunschweig.
- Rogers, Hartley** 1987 Theory of Recursive Functions and Effective Computability; 2. Auflage, MIT Press, Cambridge, MA. (Erste Auflage: Mc Graw-Hill, New York, 1967.)
- Smoryński, Craig** 1991 Logical Number Theory; Springer, Berlin.
- Shoenfield, Joseph** 1967 Mathematical Logic; Addison-Wesley, Reading, Mass. (Nachdruck 2001 bei Peters, Natick, MA.)
- Rautenberg, Wolfgang** 1996 Einführung in die mathematische Logik; Vieweg, Braunschweig. (Englische Übersetzung "A Concise Introduction to Mathematical Logic", 2. Auflage, Springer 2006.)
- Schütte, Kurt** 1960 Beweistheorie; Springer, Berlin.
- Soare, Robert** 1987 Recursively Enumerable Sets and Degrees; Springer, Berlin.
- Tarski, Alfred** 1977 Einführung in die mathematische Logik; 5. Auflage. (Übersetzung aus dem Englischen von Erhard Scheibe.) Vandenhoeck & Ruprecht, Göttingen.
- Troelstra, Anne Sierp / Schwichtenberg, Helmut** 2000 Basic Proof Theory; zweite Auflage (erste Auflage 1996), Cambridge University Press, Cambridge.

11. Historische Arbeiten

- Becker, Oskar** 1964 Grundlagen der Mathematik in geschichtlicher Entwicklung; *zweite Auflage*, Karl Alber, Freiburg.
- 1973 Mathematische Existenz. Untersuchungen zur Logik und Ontologie mathematischer Phänomene; *Nachdruck der ersten Auflage 1927*. Max Niemeyer, Tübingen.
- Berka, Karel / Kreiser, Lothar** 1986 Logik-Texte. Kommentierte Auswahl zur Geschichte der modernen Logik; *Vierte Auflage*. Akademie-Verlag, Berlin.
- Bochenski, Józef Maria** 1978 Formale Logik; *Unveränderter Neudruck der zweiten erweiterten Auflage*. Alber, Freiburg.
- Brieskorn, Egbert** (Hrsg.) 1996 Felix Hausdorff zum Gedächtnis; *Vieweg, Braunschweig*.
- van Dalen, Dirk / Monna, A. F.** 1972 Sets and Integration; *Groningen*.
- van Dalen, Dirk / Ebbinghaus, Heinz-Dieter** 2000 Zermelo and the Skolem Paradox; *Bulletin of Symbolic Logic* 6 (2000), S. 145 – 161.
- Dauben, Joseph Warren** 1971 The trigonometric background to Georg Cantor's theory of sets; *Archive for History of Exact Sciences* 7, S. 181 – 216.
- 1979a Georg Cantor's creation of transfinite set theory: personality and psychology in the history of mathematics; *Annals of the New York Academy of Sciences* 321 (1979), S. 27 – 44.
 - 1979b Georg Cantor – His Mathematics and Philosophy of the Infinite; *Harvard University Press, Cambridge, Mass.* (Nachdruck 1990, Princeton University Press.)
 - 1981 (Hrsg.) Mathematical Perspectives. Essays on Mathematics and its Historical Development; *Academic Press, New York*.
- Deiser, Oliver** 2005 Der Multiplikationssatz der Mengenlehre; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 107 (2005), S. 89 – 109.
- 2008 Essay zu Zermelos Arbeit „Ueber die Addition transfiniter Cardinalzahlen“. in Band I der *Gesammelten Werke von Ernst Zermelo* (Hrsg. Heinz-Dieter Ebbinghaus, Akihiro Kanamori, Craig Fraser). Springer, Berlin.
- Dierkesmann, Magda et al.** 1967 Felix Hausdorff zum Gedächtnis; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 69, S. 51 – 76.
- Dugac, Pierre** 1976 Richard Dedekind et les fondements des mathématiques (avec de nombreux textes inédits); *Vrin, Paris*.
- Ebbinghaus, Heinz-Dieter** 2007 Ernst Zermelo. An approach to his life and work; *In Zusammenarbeit mit Volker Peckhaus*. Springer, Berlin.
- Eichhorn, Eugen / Thiele, Ernst-Jochen** (Hrsg.) 1994 Vorlesungen zum Gedenken an Felix Hausdorff; *Berliner Studienreihe zur Mathematik* 5, Heldermann, Berlin.
- Felgner, Ulrich** 2002 Der Begriff der Funktion; in: *Felix Hausdorff: Gesammelte Werke, Band II*, S. 621 – 634. Springer, Berlin.
- Ferreirós, José** 1996 Traditional Logic and the early history of sets, 1854 – 1908; *Archive for History of Exact Sciences* 50 (1996), S. 5 – 71.
- 1999 Labyrinths of Thought. A History of Set Theory and its Role in Modern Mathematics; *Science Networks. Historical Studies, Volume 23*. Birkhäuser, Basel.
- Fraenkel, Abraham** 1930 Georg Cantor; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 39 (1930), S. 189 – 266.

- 1967 Lebenskreise – Aus den Erinnerungen eines jüdischen Mathematikers; *Deutsche Verlags-Anstalt, Stuttgart*.
- Führich, Arnold** 1983 Der Meinungsstreit zwischen Georg Cantor und Leopold Kronecker um Grundlagen der Mathematik in der Zeit der Begründung der Mengenlehre; *Potsdamer Forschungen 34, Pädagogische Hochschule „Karl Liebknecht“, Potsdam*.
- Garciadiego, Alejandro** 1992 Bertrand Russell and the Origins of the Set-Theoretic “Paradoxes”; *Birkhäuser, Basel*.
- Gericke, Helmuth** 1970 Geschichte des Zahlbegriffs; *Bibliographisches Institut, Mannheim*.
- 1971 Zur Geschichte des Zahlbegriffs; *Mathematisch-Physikalische Semesterberichte 18 (1971), S. 161 – 173*.
- 1973 Vorgeschichte der Mengenlehre; *Mathematisch-Physikalische Semesterberichte 20 (1973), S. 151 – 170*.
- 1977 Wie vergleicht man unendliche Mengen?; *Sudhoffs Archiv 61 (1977), S. 54 – 65*.
- 1980 Wie dachten und wie denken die Mathematiker über das Unendliche?; *Karl-Sudhoff-Gedächtnisvorlesung 1979. Sudhoffs Archiv 64 (1980), S. 207 – 225*.
- Gottwald, Siegfried / Kreiser, Lothar** 1984 Paul Mahlo. Leben und Werk; *Schriftenreihe für Geschichte der Naturwissenschaften, Technik und Medizin 21 (1984), S. 1 – 22*.
- Gottwald, Siegfried / Ilgauds, Hans-Joachim / Schlote, Karl-Heinz** 1990 Lexikon bedeutender Mathematiker; *Harri Deutsch, Frankfurt*.
- Grattan-Guinness, Ivor** 1970 An unpublished paper by Georg Cantor: Prinzipien einer Theorie der Ordnungstypen. Erste Mitteilung; *Acta Mathematica 124 (1970), S. 65 – 107*.
- 2000 (Hrsg.) From the Calculus to Set Theory, 1630 – 1910; *Princeton University Press, Princeton, NJ. (Nachdruck der 1. Auflage 1980, Duckworth, London.)*
- 2000 The Search for Mathematical Roots, 1870 – 1930. Logic, Set Theories and The Foundations of Mathematics from Cantor Through Russell to Gödel; *Princeton University Press, Princeton, NJ*.
- Hallett, Michael** 1984 Cantorian set theory and limitation of size; *Clarendon Press, Oxford*.
- Hintikka, Jaakko** (Hrsg.) 1995 From Dedekind to Gödel; *Dordrecht, Kluwer*.
- Ilgauds, Hans-Joachim** 1985 Zur Biographie von Felix Hausdorff; *Mitteilungen der Mathematischen Gesellschaft der DDR, Heft 2-3 (1985), S. 59 – 70*.
- Jourdain, Philip** 1906 – 1914 The development of the theory of transfinite numbers; *Archiv für Mathematik und Physik (Grunerts Archiv) 10 (1906), S. 254 – 281; 14 (1908/1909), S. 287 – 311; 16 (1910), S. 21 – 43; 22 (1913/1914), S. 1 – 21*.
- 1910 – 1913 The development of theories of mathematical logic and the principles of mathematics; *Quarterly Journal for pure and applied mathematics 41 (1910), S. 324 – 352; 43 (1912), S. 219 – 314; 44 (1913), S. 113 – 128*.
- 1991 Selected essays on the history of set theory and logics (1906 – 1918). Edited by Ivor Grattan-Guinness; *Sources for the History of Logic in the Modern Age, VI. Cooperativa Libreria Universitaria Editrice Bologna, Bologna*.
- Kanamori, Akihiro** 1996 The mathematical development of set theory from Cantor to Cohen; *The Bulletin of Symbolic Logic, Volume 2, Number 1 (1996), S. 1 – 71*.
- 2003 The empty set, the singleton, and the ordered pair; *The Bulletin of Symbolic Logic, Volume 9, Number 3 (2003), S. 273 – 298*.

- 2004 Zermelo and set theory; *The Bulletin of Symbolic Logic* 10 (2004), S. 487 – 553.
- Kertész, Andor** 1983 Georg Cantor (1845 – 1918) – Schöpfer der Mengenlehre; *bearbeitet von Manfred Stern; Acta Historica Leopoldina* 15. Barth, Leipzig.
- Maier, Anneliese** 1949 Die Vorläufer Galileis im 14. Jahrhundert; *Edizioni di storia e letteratura, Rom. (Neudruck mit Nachträgen 1966.)*
- Manheim, Jerome** 1964 The Genesis of Point Set Topology; *Macmillan, New York.*
- Meschkowski, Herbert** 1967 Probleme des Unendlichen. Werk und Leben Georg Cantors; *Vieweg, Braunschweig.*
- 1973 Hundert Jahre Mengenlehre; *Deutscher Taschenbuch Verlag, München.*
- Moore, Gregory H.** 1978 The origins of Zermelo's axiomatisation of set theory; *Journal of Philosophical Logic* 7 (1978), S. 307 – 329.
- 1980 Beyond first order logic: The historical interplay between mathematical logic and axiomatic set theory; *History and Philosophy of Logic* 1 (1980), S. 95 – 137.
- 1982 Zermelo's Axiom of Choice: Its Origins, Development, and Influence; *Springer, New York.*
- Nidditch, Peter H.** 1987 The Development of Mathematical Logic; *Routledge, London.*
- Peckhaus, Volker** 1990 Hilbertprogramm und Kritische Philosophie; *Vandenboeck & Ruprecht, Göttingen.*
- 1990b „Ich habe mich wohl gehütet, alle Patronen auf einmal zu verschießen.“ Ernst Zermelo in Göttingen; *History and Philosophy of Logic* 11, S. 19 – 58.
- 1997 Zermelo in Zürich; *Vortrag, gehalten am 15.7.1997 in Göttingen.*
- Purkert, Walter** 1986 Georg Cantor und die Antinomien der Mengenlehre; *Bulletin de la Société Mathématique de Belgique* 38 (1986), S. 313 – 327.
- 2002 Grundzüge der Mengenlehre – Historische Einführung; *in: Felix Hausdorff: Gesammelte Werke, Band II, S. 1 – 91. Springer, Berlin.*
- Purkert, Walter / Ilgauds, Hans Joachim** 1987 Georg Cantor 1845 – 1918; *Birkhäuser Verlag, Basel. Überarbeitete Fassung der Ausgabe 1985 bei Teubner, Leipzig.*
- Rautenberg, Wolfgang** 1987 Über den Cantor-Bernsteinschen Äquivalenzsatz; *Mathematisch-Physikalische Semesterberichte* 34 (1987), S. 71 – 88.
- Rowe, D. / McCleary, J.** (Hrsg.) 1989 The History of Modern Mathematics. Volume I: Ideas and their Reception; *Academic Press, Boston.*
- Scharlau, Winfried** (Hrsg.) 1981 Richard Dedekind. Eine Würdigung zu seinem 150. Geburtstag; *Vieweg, Braunschweig.*
- Schoenflies, Arthur** 1922 Zur Erinnerung an Georg Cantor; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 31 (1922), S. 97 – 106.
- 1927 Die Krisis in Cantors mathematischem Schaffen; *Acta Mathematica* 50 (1927), S. 1 – 23.
- Simons, Peter** 1992 Philosophy and Logic in Central Europe from Bolzano to Tarski. With a presentation by Witold Marciszewski. Selected Essays; *Nijhoff International Philosophy Series, 45. Kluwer Academic Publishers Group, Dordrecht.*
- Styazhkin, Nikolai Ivanovich** 1969 History of Mathematical Logic from Leibniz to Peano; *Cambridge, Mass.*
- Tapp, Christian** 2005 Kardinalität und Kardinäle : Wissenschaftshistorische Aufarbeitung der Korrespondenz zwischen G. Cantor und kath. Theologen seiner Zeit; *Steiner, Stuttgart.*

- Thiel, Christian** 1984 Folgen der Emigration deutscher und österreichischer Wissenschaftstheoretiker und Logiker zwischen 1933 und 1945; *Berichte zur Wissenschaftsgeschichte* 7 (1984), S. 227 – 256.
- Tsouyopoulos, Nelly** 1972 Der Begriff des Unendlichen von Zenon bis Galilei; *Rete, Strukturgeschichte der Naturwissenschaften* 1 (1972), S. 245 – 272.

12. Philosophische Schriften und Anthologien

- Balaguer, Marc** 1998 Platonism and Anti-Platonism in Mathematics; *Oxford University Press, New York*.
- Beth, Evert W.** 1968 The Foundations of Mathematics; *North-Holland, Amsterdam*.
- Benacerraf, Paul / Putnam, Hilary** (Hrsg.) 1983 Philosophy of Mathematics; 2. Auflage. *Cambridge University Press, Cambridge*.
- Bernays, Paul** 1976 Abhandlungen zur Philosophie der Mathematik; *Wissenschaftliche Buchgesellschaft, Darmstadt*.
- Bolzano, Bernard** 1851 Paradoxien des Unendlichen; *Reclam, Leipzig*; *Reprographischer Nachdruck bei Felix Meiner, Hamburg* 1955.
- Cavaillès, Jean** 1962 Philosophie mathématique; *Hermann, Paris*.
- Dales, H. Garth / Olivieri, Gianluigi** (Hrsg.) 1998 Truth in Mathematics; *Clarendon Press, Oxford*.
- Erickson, Glenn / Fossa, John** 1998 Dictionary of Paradox; *University Press of America, Lanham*.
- Finsler, Paul** 1925 Gibt es Widersprüche in der Mathematik?; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 34 (1925), S. 143 – 155.
- 1975 Aufsätze zur Mengenlehre; Hrsg. Georg Unger. Mit einem Geleitwort und einem Nachruf auf Finsler von Herbert Gross. *Wissenschaftliche Buchgesellschaft, Darmstadt*.
- Fraenkel, Abraham / Bar-Hillel, Yehoshua / Levy, Azriel** 1973 Foundations of Set Theory; 2. Auflage, *North-Holland, Amsterdam*.
- Galilei, Galileo** 1638 Discorsi e dimostrazioni matematiche; *Leiden, 1638*.
- George, Alexander / Velleman, Daniel** 2002 Philosophies of Mathematics; *Blackwell, Malden, Mass.*
- Gödel, Kurt** 1947 What is Cantor's Continuum Problem?; *American Mathematical Monthly* 54 (1947), S. 515 – 525.
- Hart, W.D.** (Hrsg.) 1996 The Philosophy of Mathematics; *Oxford University Press, Oxford*.
- Hessenberg, Gerhard** 1908 Willkürliche Schöpfungen des Verstandes?; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 17, S. 145 – 162.
- Hilbert, David** 1925 Über das Unendliche; *Mathematische Annalen* 95 (1925), S. 161 – 190.
- 1964 Hilbertiana – Fünf Aufsätze von David Hilbert; *Wissenschaftliche Buchgesellschaft, Darmstadt*. (Enthält: Axiomatisches Denken, Neubegründung der Mathematik, Die logischen Grundlagen der Mathematik, Die Grundlagen der Physik, Über das Unendliche.)
- Hilbert, David / Bernays Paul** 1968, 1970 Grundlagen der Mathematik; 2 Bände, 2. Auflage, *Springer, Berlin*.
- Jacquette, Dale** (Hrsg.) 2002 a Philosophy of Logic. An Anthology; *Blackwell, Oxford*.

- 2002 b (Hrsg.) *Philosophy of Mathematics. An Anthology*; *Blackwell, Oxford*.
- Korselt, Alwin** 1911 Über mathematische Erkenntnis; *Jahresbericht der Deutschen Mathematiker-Vereinigung* 20 (1911), S. 364 – 380.
- Lavine, Shaughan** 1994 *Understanding the Infinite*; *Harvard University Press, Cambridge*.
- Link, Godehard** (Hrsg.) 2004 *One Hundred Years of Russell's Paradox*; *de Gruyter, Berlin*.
- Maddy, Penelope** 1990 *Realism in Mathematics*; *Clarendon Press, Oxford*.
- 1997 *Naturalism in Mathematics*; *Clarendon Press, Oxford*.
- Moravcsik, Julius** 1992 *Plato and Platonism. Plato's conception of appearance and reality in ontology, epistemology and ethics, and its modern echoes.*; *Blackwell, Oxford*.
- Nelson, Leonard** 1959 *Beiträge zur Philosophie der Logik und Mathematik*; *mit Einführungen von W. Ackermann, P. Bernays, D. Hilbert. Verlag öffentliches Leben, Frankfurt*.
- Parsons, Charles** 2005 *Mathematics in Philosophy*; *Cornell University Press, Cornell*.
- Russell, Bertrand** 2002 *Einführung in die mathematische Philosophie*; *Übersetzung der englischen Ausgabe von 1919 durch Emil Gumbel, bearbeitet von Johannes Lengard, Michael Otte. Meiner, Hamburg*.
- Schirn, Matthias** (Hrsg.) 1998 *The Philosophy of Mathematics Today*; *Clarendon Press, Oxford*.
- Shanker, Stuart** (Hrsg.) 2001 *Philosophy of Science, Logic and Mathematics in the Twentieth Century*; *Nachdruck der Originalausgabe von 1996. Routledge London*.
- Tymoczko, Thomas** (Hrsg.) 1998 *New Directions in the Philosophy of Mathematics. An Anthology. Revised and expanded edition*; *Princeton University Press, Princeton*.
- Thiel, Christian** (Hrsg.) 1982 *Erkenntnistheoretische Grundlagen der Mathematik*; *Gerstenberg, Hildesheim*.
- 1995 *Philosophie und Mathematik. Eine Einführung in ihre Wechselwirkungen und in die Philosophie der Mathematik*; *Wissenschaftliche Buchgesellschaft, Darmstadt*.
- Tiles, Mary** 1989 *The Philosophy of Set Theory. An Introduction to Cantor's Paradise*; *Blackwell, Oxford*.
- Wang, Hao** 1974 *From Mathematics to Philosophy*; *Routledge, London*.
- Wedberg, Anders** 1955 *Plato's Philosophy of Mathematics*; *Almqvist & Wiksell, Stockholm*.

13. Nichtmathematische Schriften von Georg Cantor

- Cantor, Georg** (Hrsg.) 1896 a *Confessio fidei Francisci Baconi Baronis de Verulam vicecomitis Sancti Albani. Anglico sermone ante annum MDCIV conscripta; cum versione Latina a Guilelmo Rawley, s. theol. doctore dominationi suae a sacris et operum ejus editore, anno MDCLVIII evulgata. – Nunc denuo typis excusa cura et impensis G. C. – Halis Saxonum MDCCCXCVI. – Prostat apud Max. Niemeyer bibliopolam; Mit lateinischer Vorrede von Georg Cantor. Niemeyer, Halle*.
- 1896 b (Hrsg.) *Resurrectio Divi Quirini Francisci Baconi Baronis de Verulam vicecomitis Sancti Albani: CCLXX annis post obitum eius IX die Aprilis anni MDCXXVI. – (Pro manuscripto.) – Cura et impensis G. C. – Halis Saxonum MDCCCXCVI; mit engl. Vorrede von „Dr. phil. George Cantor, Mathematicus.“ Selbstverlag*.
- 1897 (Hrsg.) *Die Rawley'sche Sammlung von zweiunddreißig Trauergedichten auf Francis Bacon: ein Zeugnis zugunsten der Bacon-Shakespeare-Theorie; mit einem*

Vorwort herausgegeben von Georg Cantor; *Nachdruck der Ausgabe London 1626. Niemeyer, Halle.*

- 1900 Shaxpeareologie und Baconianismus; *Magazin für Literatur* 69 (1900), Spalten 196 – 203.
- 1905 EX ORIENTE LUX. Gespräche eines Meisters mit seinem Schüler über wesentliche Punkte des urkundlichen Christentums. Berichtet vom Schüler selbst Georg Jacob Aaron, cand. sacr. theol. Erstes Gespräch. Pro manuscripto. Herausgegeben von Georg Cantor. Im Selbstverlag des Herausgebers; *Halle.*

14. Schriften von Felix Hausdorff alias Paul Mongré

- Mongré, Paul** 1897a Sant' Ilario – Gedanken aus der Landschaft Zarathustras; *Naumann, Leipzig.*
- 1897b Selbstanzeige: Sant' Ilario – Gedanken aus der Landschaft Zarathustras; *Die Zukunft*, 20.11.1897, S. 361.
 - 1898a Das Chaos in kosmischer Auslese – Ein erkenntniskritischer Versuch; *Naumann, Leipzig.*
 - 1898b Massenglück und Einzelglück; *Neue Deutsche Rundschau* 9 (1) (1898), S. 64 – 75.
 - 1898c Das unreinliche Jahrhundert; *Neue Deutsche Rundschau* 9 (5) (1898), S. 443 – 452.
 - 1898c Stirner; *Die Zeit* 213, 29.10.1898, S. 69 – 72.
 - 1899a Tod und Wiederkunft; *Neue Deutsche Rundschau* 10 (12) (1899), S. 1277 – 1289.
 - 1899b Selbstanzeige: Das Chaos in kosmischer Auslese; *Die Zukunft* 8 (5) (1899), S. 222 – 223.
 - 1900 Ekstasen; *Gedichte. Seemann, Leipzig.*
 - 1900a Nietzsches Wiederkunft des Gleichen; *Die Zeit* 292, 5.5.1900, S. 72 – 73.
 - 1900b Nietzsches Lehre von der Wiederkunft des Gleichen; *Die Zeit* 297, 9.6.1900, S. 150 – 152.
 - 1902a Der Schleier der Maja; *Neue Deutsche Rundschau* 13 (9) (1902), S. 985 – 996.
 - 1902b Der Wille zur Macht; *Neue Deutsche Rundschau* 13 (12) (1902), S. 1334 – 1338. (Rezension von „Der Wille zur Macht“, einer Kompilation aus dem Nachlaß von Nietzsche.)
 - 1902c Max Klingers Beethoven; *Zeitschrift für bildende Kunst, Neue Folge* 13 (1902), S. 183 – 189.
 - 1902d Brief gegen G. Landauers Artikel „Die Welt als Zeit“; *Die Zukunft* 10 (37) (1902), 14.6.1902, S. 441 – 445.
 - 1903 Sprachkritik; *Neue Deutsche Rundschau* 14 (12) (1903), S. 1233 – 1258.
 - 1904a Gottes Schatten; *Die neue Rundschau* 15 (1) (1904), S. 122 – 124.
 - 1904b Der Arzt seiner Ehre. Groteske; *Die neue Rundschau* 15 (8) (1904), S. 989 – 1013. (Buchausgaben: *Leipziger Bibliophilen-Abend*, 1910 und S. Fischer, Berlin, 1912.)
 - 1909 Strindbergs Blaubuch; *Die neue Rundschau* 20 (6) (1909), S. 891 – 896.
 - 1910a Der Komet; *Die neue Rundschau* 21 (5) (1910), S. 708 – 712.
 - 1910b Andacht zum Leben; *Die neue Rundschau* 21 (12) (1910), S. 1737 – 1741.
 - 1912 Biologisches; *Licht und Schatten* 3 (1912/1913), unpaginiert.

6. Notationen

Mengenbildung

$\{a_1, \dots, a_n\}$, 30, 423
 $\mathcal{E}(x)$, 32
 $\{x \mid \mathcal{E}(x)\}$, 32
 $\{\mathcal{F}(x) \mid \mathcal{E}(x)\}$, 39

Spezielle Mengen

$\mathbb{N}$, 27
 $\mathbb{Z}$, 27
 $\mathbb{Q}$, 27
 $\mathbb{R}$, 27
 $\mathbb{A}$, 117
 $\mathbb{T}$, 122
 $\mathfrak{S}$, 139
 $\emptyset$, 31, 422
 $\{ \}$, 31
 $\bar{n}$, 100
 $\mathcal{P}^*(\mathbb{N})$, 134
 C , 376

Spezielle Klassen

V , 183
 R , 183
 Ω , 401
 $Kard$, 402
 $\beth$, 402

Relationen

$a \in b$, 15
 $a \subseteq b$, 33
 $a \subset b$, 33
 $b \supseteq a$, 33
 $b \supset a$, 33
 $a R b$, 51

Operationen

$a \cup b$, 35, 423
 $a \cap b$, 35, 426
 $a - b$, 35, 426
 a^c , 36
 $a \Delta b$, 38
 $\bigcap M$, 41, 426
 $\bigcup M$, 41, 423
 $\mathcal{P}(M)$, 41, 428
 (a, b) , 48, 423
 $A \times B$, 50, 428
 $(a_1, \dots, a_n)$, 50
 $A_1 \times \dots \times A_n$, 50
 $\text{dom}(R)$, 51
 $\text{rng}(R)$, 51
 A/R , 52
 R^{-1} , 62
 $R \circ S$, 62
 $R''A$, 62
 $R^{-1}''B$, 62
 $2M$, 136

Funktionen

$f(a) = b$, 54
 $f : A \rightarrow B$, 55
 $f \upharpoonright C$, 58
 f^{-1} , 58
 $g \circ f$, 59
 $f''A$, 61
 $f[A]$, 61
 ${}^A B$, 133

Spezielle Funktionen

$\pi : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$, 112
 $\rho : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$, 113
 id_A , 59
 $\text{ind}_{A,M}$, 134

Reelle Zahlen und Punktmengen

$\min$, 27, 29
 $|x|$, 29
 $X \leq a$, 123
 $\sup(X)$, 123
 $]x, y[$, 124, 208
 $\sup(P)$, 207
 $\inf(P)$, 207
 $U_\varepsilon(a)$, 208
 $\lim_n \rightarrow_\infty x_n$, 210
 P' , 212
 $\text{cl}(P)$, 215
 $\text{int}(P)$, 215
 $P^{(n)}$, 216
 $P^{(\omega)}$, 217
 $P^{(\alpha)}$, 360
 P^* , 364
 $\text{cp}(P)$, 367
 $\text{Pr}(P)$, 396
 Σ_1^1 , 397
 Π_1^1 , 397
 $\mathcal{G}$, 279
 $\mathcal{F}$, 279
 Δ_α , 279

Mächtigkeiten

$|A| = |B|$, 67
 $|A| \leq |B|$, 70
 $|A| \geq^* |B|$, 71
 $|A| < |B|$, 77
 $|M|$, 162

Kardinalzahlen

α, b, c , 163
 n , 163
 $\aleph_0$, 163
 c , 164
 $\mathfrak{f}$, 165
 $|M|$, 265
 $|\alpha|$, 265
 ω_1 , 266
 $\aleph_1$, 266
 $\beth_\alpha$, 275
 $\aleph_\alpha$, 278
 $\aleph(\kappa)$, 310

**Kardinalzahlen,
Arithmetik**

$\alpha + b$, 164
 $\alpha \cdot b$, 164
 α^b , 164
 α^+ , 171
 $[A]^b$, 177
 $\min(\mathfrak{A})$, 180
 $\kappa + \lambda$, 265
 $\kappa \cdot \lambda$, 265
 κ^λ , 265
 κ^+ , 266

**Spezielle
Ordnungstypen**

n , 285
 ω , 285
 ζ , 285
 η , 285, 341
 θ , 285, 341
 ε_0 , 294
 ξ , 379

Ordnungstheorie

ω , 203
 $X < s$, 224
 $X \leq s$, 224
 $<|N$, 224
 $\bigcup \Gamma$, 227
 $\equiv$, 230, 284
 $<|$, 232
 $\mathcal{H}(M)$, 242
 $a \parallel b$, 244
 $\text{o.t.}(\langle M, < \rangle)$, 252
 $W(\alpha)$, 253
 $\sup(A)$, 255
 $\sup(A)$, 260
 o.t. , 262
 α^* , 286
 $W(\alpha, \beta)$, 297
 Γ , 297
 $\text{cf}(\langle M, < \rangle)$, 306
 $\text{cf}(\alpha)$, 306
 $\text{Lim}(X)$, 320
 $\sup$, 345
 $\inf$, 345
 $\alpha \leq \beta$, 349
 $\alpha \leq^* \beta$, 349
 $\mathcal{D}(M)$, 354

**Ordnungstheorie,
Arithmetik**

$\langle M, < \rangle + \{x\}$, 228
 $\alpha + 1$, 254, 260
 $\alpha - 1$, 254, 260
 $\langle M, < \rangle + \langle N, < \rangle$, 287
 $\alpha + \beta$, 287
 $\langle M, < \rangle \cdot \langle N, < \rangle$, 288
 $\alpha \cdot \beta$, 289
 $\text{exp}_{\text{lex}}(\langle M, < \rangle, \langle N, < \rangle)$, 292
 $\text{exp}_{\text{lex}}(\alpha, \beta)$, 292
 α^β , 293
 $\langle M, < \rangle^{(N, <)}$, 302

Logik

$\mathcal{L}$, 449
 $\neg, \wedge, \vee, \rightarrow, \leftrightarrow, \forall, \exists$, 450
 $\cdot$, 451
 $\varphi(x_1, \dots, x_n)$, 452
 $\exists !$, 454
 $\vdash$, 460

Klassennotationen

$\{x \mid \varphi(x, p_1, \dots, p_n)\}$, 469
 V , 469
 $A = B$, 469
 $A \in B$, 469
 $A \subseteq B$, 469
 $A \cup B$, 470
 $A \cap B$, 470
 $A - B$, 470
 $\bigcup A$, 470
 $\bigcap A$, 470
 $A \times B$, 470
 $\mathcal{P}(A)$, 470
 $F(x) = y$, 471
 $\in$, 471
 $F''X$, 472
 $F \upharpoonright X$, 472
 $G \circ F$, 472
 id , 472
 $R \circ_{\text{Rel}} S$, 472

**Abkürzungen und
Bezeichnungen**

gdw , 30
 CH , 149, 266
 GCH , 157, 266
 L , 151, 328
 V_α , 276
 SCH , 310
 ZFC , 417
 NBG , 440, 475
 MK , 440, 477
 Z, ZF , 441
 $\text{ZF}^-, \text{ZFC}^-$, 441

Hebräisches Alphabet		Griechisches Alphabet			Fraktur Alphabet		
Aleph	א	Alpha	A	α	A	ℒ	α
Beth	ב	Beta	B	β	B	℔	β
Gimel	ג	Gamma	Γ	γ	C	℥	γ
Daleth	ד	Delta	Δ	δ	D	℄	δ
He	ה	Epsilon	E	ε	E	℥	ε
Waw	ו	Zeta	Z	ζ	F	ℷ	Ϝ
Zayin	ז	Eta	H	η	G	™	g
Heth	ח	Theta	Θ	θ, ϑ	H	℥	h
Teth	ט	Jota	I	ι	I	ℑ	i
Yodh	י	Kappa	K	κ	J	ℑ	j
Kaph	כ	Lambda	Λ	λ	K	℔	k
Lamedh	ל	My	M	μ	L	℔	l
Mem	מ	Ny	N	ν	M	℞	m
Nun	נ	Xi	Ξ	ξ	N	℔	n
Samekh	ס	Omikron	O	ο	O	℄	o
Ayin	ע	Pi	Π	π	P	℔	p
Pe	פ	Rho	P	ρ	Q	℄	q
Sadhe	צ	Sigma	Σ	σ, ς	R	℔	r
Qoph	ק	Tau	T	τ	S	℥	s
Resh	ר	Ypsilon	Y	υ	T	℥	t
Shin	ש	Phi	Φ	φ	U	ℒ	u
Taw	ת	Chi	X	χ	V	℔	v
		Psi	Ψ	ψ	W	℔	w
		Omega	Ω	ω	X	℔	x
					Y	℔	y
					Z	℔	z

Skriptbuchstaben

A	ℒ	J	ℑ	S	ℑ
B	℔	K	℔	T	℥
C	℥	L	℔	U	ℒ
D	℄	M	℔	V	℔
E	℥	N	℔	W	℔
F	ℷ	O	℣	X	℔
G	™	P	℔	Y	℔
H	℥	Q	℔	Z	℔
I	ℑ	R	℔		

7. Personenverzeichnis

Albert von Sachsen, 68
Alexandrov, Pavel, 325, 395ff, 412, 486,
495f, 503ff
Anaximander, 68
Aristoteles, 68f, 359, 456

Bacon, Roger, 68
Baire, René, 137ff, 381, 396, 399, 468
486, 505
Banach, Stefan, 244, 336, 440, 495, 511
Bendixson, Ivar, 3, 201, 214, 360ff, 375f,
382ff, 389, 392, 486, 489, 489, 500
Bernays, Paul, 440, 474f, 504
Bernstein, Felix, 5, 42, 65, 71ff, 87f, 96,
98, 114, 125, 128ff, 137, 150, 166,
175, 198, 236, 243, 246, 285, 394f,
395, 399f, 400, 421, 439, 479f, 486,
491f, 492, 497, 501f, 509
Berry, G.G. 191f
Bloch, Gérard, 324
BoisReymond, Paul Du, 25
Bolzano, Bernard, 19, 69, 91, 209ff, 486,
499
Borel, Émile, 339, 480, 486, 489, 501
Brouwer, Luitzen, 131, 196, 458
Bruno, Giordano, 116
Bruns, Heinrich, 407
Burali-Forti, Cesare, 10, 89. 191, 401
Buridan, 68

Cantor, Georg, *passim*
Cauchy, Augustin-Louis, 24, 211, 365,
487
Cohen, Paul, 11, 152, 156f, 358, 486,
499, 506f, 510

Dedekind, Richard, 1, 16f, 20f, 42, 55,
74, 76, 89, 93ff, 98ff, 118, 120ff, 130f,
144, 152, 163, 189, 196ff, 246, 345,
354ff, 358, 391, 401ff, 418, 420, 430,
440, 447, 465, 467, 486, 488, 491, 497,
499, 501, 509, 512
Dirichlet, Peter Gustav Lejeune, 55, 196

Einstein, Albert, 108, 199
Euklid, 20, 68, 108, 264, 495

Euler, Leonhard, 55, 122

Fermat, Pierre de, 31
Fourier, Jean-Baptiste, 55, 69, 407, 412
Fraenkel, Abraham, 6, 10, 24f, 46f, 57, 67,
79f, 161, 193ff, 234, 246, 263, 415, 417,
420f, 433, 435, 437, 477, 480f, 486, 497f,
502ff, 506, 510
Fréchet, Maurice, 343, 411, 494
Frege, Gottlob, 184, 456, 509

Galilei, Galileo, 69
Gauß, Carl Friedrich, 24, 112, 144, 196
Gelfond, Alexander, 122
Gentzen, Gerhard, 456
Gödel, Kurt, 11, 46, 152, 156f, 358, 418, 438,
440, 442, 474, 486
Grelling, Kurt, 191f

Hadamard, Jacques, 198
Hankel, Hermann, 25
Harnack, Axel, 25
Hartogs, Friedrich, 240ff, 248, 275, 281, 402,
486, 503, 509
Hausdorff, Felix, *passim*
Helmholtz, Hermann von, 144
Herbart, Johann Friedrich, 23, 404
Hermite, Charles, 122
Hessenberg, Gerhard, 41, 66, 94, 141, 175f,
179, 298, 301, 303ff, 458, 486, 490, 492,
502, 510
Heyting, Arend, 458
Hilbert, David, 10f, 24, 45f, 89, 91f, 107f,
122, 142, 147, 184, 194, 197ff, 249, 401,
404, 411, 417, 424, 442f, 453, 456ff, 479f,
486, 496, 501
Hurwitz, Adolf, 198

Jourdain, Philip, 176, 298, 480, 492, 502, 512

Kolmogorov, Andrej, 412, 458
König, Julius, 8, 75, 130, 132, 167ff, 178,
191, 198, 309, 479f, 486
Korselt, Alwin, 74
Kronecker, Leopold, 17, 195ff, 480
Kummer, Ernst Eduard, 195

Kuratowski, Kazimierz, 48ff, 156, 314, 336, 486, 491f, 510

Lebesgue, Henri, 53, 214, 336, 339, 382, 398f, 440, 486, 489, 495, 505

Leibniz, Gottfried Wilhelm, 20, 55

Lindemann, Ferdinand von, 113, 122

Liouville, Joseph, 119, 122

Lipschitz, Rudolf, 196

Lusin, Nikolai, 397, 399, 412, 486, 496, 503, 505f, 509f

Mahlo, Paul, 8, 306, 317, 319, 324ff, 332f, 339f, 404, 477, 486, 503, 505f, 509

Méray, Charles, 25, 356

Minkowski, Hermann, 198

Mirimanov, Dimitry, 185, 190, 276, 498

Mittag-Leffler, Gösta, 197, 383, 392

Neumann, John von, 10, 190, 251, 257, 261ff, 276ff, 282f, 403, 417, 420, 430ff, 436, 440, 474ff, 486, 492, 498, 504, 509

Nikomachos, 20, 264

Ockham, Wilhelm von, 68

Oresme, Nikolaus von, 68

Peano, Giuseppe, 21, 33, 35, 51, 74, 107, 142, 152, 465, 480, 486, 489, 495

Planck, Max, 107, 479

Platon, 21f, 118, 513

Poincaré, Henri, 74, 392

Riemann, Bernhard, 55, 141, 144, 149, 196, 499, 512

Russell, Bertrand, 10, 46, 55, 89, 160, 183ff, 191, 199, 401, 418, 426, 438, 470, 476, 486, 491, 497, 512

Scheeffer, Ludwig, 392ff, 398ff, 486, 496, 503, 505

Schmidt, Erhard, 81, 84f, 173, 249, 479ff

Schneider, Theodor, 122

Schoenflies, Arthur, 132, 243, 343, 400, 480, 486, 501, 510

Schröder, Ernst, 33, 51, 445

Schumacher, Heinrich, 24

Schwarz, Hermann, 93, 195, 479

Sierpiński, Wacław, 397, 505

Skolem, Thoralf, 10, 417, 431, 444ff, 504, 509

Steiner, Jacob, 162

Suslin, Mikhail, 157, 357f, 396ff, 486, 496, 503, 505f, 506

Tarski, Alfred, 88, 103ff, 154, 157, 193, 243, 314, 330, 332, 335, 440, 486, 495, 504f, 509, 512

Thales, 20, 264

Ulam, Stanisław, 306, 326f, 330f, 334f, 337f, 340, 486, 504f

Veblen, Oswald, 312

Vitali, Giuseppe, 53, 336

Weierstraß, Karl, 24, 107, 195, 197, 209ff, 488

Weil, André, 31

Young, William, 395

Zermelo, Ernst, 5ff, 10, 72, 74, 76, 81, 87f, 183f, 187, 198, 228, 233, 238, 240, 249, 257, 264, 276, 282, 293, 296, 353f, 406, ab 415 *passim*

Zorn, Max, 84f, 175, 244f, 435, 440

8. Sachverzeichnis

Abbildung, 57

abgeschlossene Teilmenge von $\mathbb{R}$, 213
abgeschlossen (Ordinalzahlen), 319
Abgeschlossenheit der Ableitung, 215
Ableitung, 212, 360, 453
Abschluß einer Punktmenge, 215
absolut unendlich, 189
absteigende Folgen, 361
abzählbar, 109
abzählbare Vereinigung, 115
abzählbarer Ordnungstyp, 285
Abzählbarkeit aller Bücher, 115
Abzählbarkeit von $\mathbb{Q}$ und $\mathbb{A}$, 116, 118
Addition von Kardinalzahlen, 164, 265
Addition zweier linearer Ordnungen, 287
Addition zweier Ordnungstypen, 287
Additivität, 333, 338
 σ -Additivität, 335
Adhärenzen, 369
ähnliche lineare Ordnungen, 284
ähnliche Wohlordnungen, 230
aktuell unendlich, 23
Aleph, 163
Alephparadoxon, 402
algebraische Zahlen, 117
Algorithmus des Abtragens, 64
Allabschluß, 451
Allklasse, 469
analytisch, 396
Anfangsstück einer Wohlordnung, 227
Antiketten-Bedingung, 357
Antinomie von Russell und Zermelo, 183, 470
Approximation der Cantormenge, 377
äquivalent, 67, 79
Äquivalenzrelation, 52
Äquivalenzsatz, 71, 74
Assoziativgesetze, 36
atomare Formel, 450
Aufzählung einer Wohlordnung, 256
Ausdruck, 449
Aussage, 451
Aussagen mit Klassen, 473
aussagenlogisches Axiom, 457
Aussonderung, 32
Aussonderungsschema, 425, 455
Auswahlaxiom, 437, 456, 475

axiom of choice, 417
axiomatische Methode, 421
Axiomensystem, 453
azyklisch, 97

Baire-Eigenschaft, 399

Baireraum, 137
Bairescher Kategoriensatz, 381
Barbier, 184
Bereich, 419
Bernstein-Menge, 395
beschränkt, 207
Beschränktheitsaxiom, 435
Betrag, 29
Beweis, 453
Beweisschema, 473
bijektiv, 56
Bild, 54, 61, 472
Binärdarstellung, 28
Bindungsstärke, 450
Block einer reellen Zahl, 132
Bolzano-Weierstraß, Satz von, 209
Borel-Hierarchie, 278

Cantor-Bendixson-Zerlegung, 364

Cantor-Bernstein, Satz von, 72, 74
Cantormenge, 136, 376
Cantorraum, 137
Cantorsche Ableitung, 212
Cantorsche Normalform, 299
Cantorsche Paarungsfunktion, 112
Cantorsches Paradoxon, 183
C-artig, 386
Cauchyfolge, 211
Charakterisierung der endlichen Mengen, 102
charakteristische Funktion, 134
club-Menge, 319
club-Filter, 321

De Morgansche Regeln, 37

Dedekind*-unendlich, 104
Dedekind**-unendlich, 104
Dedekind-Vervollständigung, 354
Deduktionstheorem, 461
Definitionsbereich, 51
Determiniertheit, 398

Dezimaldarstellung, 28
 Diagonalargument, 120
 Diagonalverfahren, 148
 dicht, 327, 341
 dicht in $\mathbb{R}$, 207
 dichte Teilmenge, 352
 Differenz, 35
 Differenzenketten, 39
 direkte Angabe der Elemente, 30
 Disjunktheit, 35
 Diskontinuum, 376
 Distributivgesetze, 37, 168
 Dualität, 37
 Dualitätsprinzip, 38
 dünn, 322

Echte Klasse, 189, 469
 echte Obermenge, 33
 echte Teilmenge, 33
 Eigenschaft, 32, 39, 453
 Eigenschaften als Operationen, 271
 Einbettung, 349
 Eindeutigkeitssatz, 259
 Einermenge, 30
 Einschränkung, 472
 Einschränkung einer Funktion, 58
 Element, 15
 Enderweiterung einer Wohlordnung, 228
 endlich, 93
 endlich axiomatisierbar, 476
 endliche Folge, 115
 endliche Kardinalität, 163
 Endlichkeit und natürliche Zahlen, 100
 Ersetzungsschema, 432, 455, 475
 Existenz der leeren Menge, 422, 455
 Exponentiation (Kardinalzahlen), 164, 265
 Exponentiation, lexikographische, 292
 Exponentiation, natürliche, 302
 Exponentiation von Ordinalzahlen, 293
 Extension, 31
 Extensionalitätsaxiom, 422, 454, 475
 Extensionalitätsprinzip, 15

Falsum, 458
 fertige Gesamtheit, 22
 Fibonacci-Zahlen, 272
 Filter, 318
 Fixpunkte, 295
 Fixpunktsatz, 76
 Folge, 115, 255
 Folge in einer Menge, 210
 Folge reeller Zahlen, 133
 forcing, 157
 formaler Beweis, 453
 Formalismus, 466
 Formel, 450

freie Variable, 451
 Fundamentalfolge, 211 (in $\mathbb{Q}$, 356)
 Fundamentalsatz der Mengenlehre, 150
 fundiert, 185
 Fundierungsaxiom, 434, 455
 Funktion, 54, 55, 471
 funktionale Eigenschaft, 432
 funktionale Klasse, 471

Ganze Zahlen, 27
 gdw, 30
 gebundene Variable, 451
 Generalisierung, 451
 geordnetes Paar, 48
 geschlossen, 82, 85
 Gmel-Funktion, 311
 gleiche Mächtigkeit, 67
 gleicher Limes, 356
 Gleichheitskriterium, 15
 Gleichheitskriterium, 34
 gleichlang, 230
 gleichmächtig, 67
 Grenzpunkt, 209
 Grenzwert, 210
 Grenzzahlen, 314
 große Kardinalzahl, 151, 314, 339, 398, 437, 465
 große Vereinigung, 41
 Größenvergleich, 64
 Größenvergleich mit Funktionen, 65
 großer Durchschnitt, 41
 Grundobjekte, 26
 Gültigkeitsrelation, 154

Hartogswohlordnung, 242
 Häufungspunkt, 209
 Hauptzahl, 299, 300
 Hausdorff-Formel, 308
 Hausdorff-Hessenberg Darstellung, 303
 Hausdorff-Residuum, 371
 Hausdorffs Maximalprinzip, 84, 247
 Hessenbergsumme, 301
 Hilbert-Kalkül, 456
 Hilbertsches Hotel, 92
 Hotelanekdoten, 481

Ideal, 319
 Identität, 59, 472
 Identitätsaxiom, 459
 in sich dicht, 213
 Indikatorfunktion, 134
 Induktion, 71, 269
 Induktionsanfang, -schritt, 71
 induktiv, 186
 induktive Definition, 271
 induzierte Wohlordnung, 231, 256
 Infimum, 207

Inhalt, 335
 injektiv, 56
 Inklusionssatz, 72
 inkompatibel, 244
 inkonsistente Vielheiten, 189
 innere Modelle, 157
 Inneres einer Punktmenge, 215
 Interpretation der Paradoxien, 187
 Intervall, 357, 361
 Intervallordnung, 386
 Intervallzerlegung, 385
 intuitionistische Logik, 458
 inverse Ordnung, 286
 isolierter Punkt, 209
 iterativ, 16
 iterierte Ableitung, 216

Junktoren, 449

Kalkül, 453
 kanonische Darstellung, 28
 kardinaler Nachfolger, 171, 266
 kardinales Supremum, 171
 Kardinalität, 162, 265
 Kardinalzahl, 78, 79, 162, 264
 kartesisches Produkt, 50, 428
 Kategorie, von erster und zweiter, 381
 Kern, 360
 Kette, 81, 85, 105, 245
 Kettenbruch, 139
 Ketten-endlich, 105
 Klasse, 189, 469
 Klassen als echte Objekte, 474
 Klassen als Relationen, 471
 klassische Logik, 458
 koanalytisch, 396
 Kodierung, 114
 Kohärenz einer Punktmenge, 369
 koinitial, 306
 komager, 380
 kompakt, 217
 Komplement, 37
 Komprehensionsprinzip, 33
 Komprehensionsschema, 475, 477
 Kondensationspunkt, 367
 konfinal, 306, 312
 Konfinalität, 306
 konsistente Vielheit, 189
 Kontinuumshypothese, 149, 266
 Kontrapositionsgesetz, 458
 konvergente Folge, 210
 korrekte Einbettung, 349
 Korrektheit (Modelle), 154
 Kreuzprodukt, 50, 167
 Kunstsprache, 444
 kürzer als, 232

Längenmaß, 382
 Lebesgue-meßbar, 399
 leere Menge, 31
 letzte Kohärenz, 371
 letztes Glied einer Kette, 105
 letztes Residuum, 373
 lexikographische Exponentiation, 292
 lexikographische Ordnung, 292
 Limes, 210, 356
 Limeselement, 226
 Limesordinalzahl, 254, 260
 Limeskardinalzahl, 267
 Limesstufe, 87
 lineare Ordnung, 224
 lineare Punktmenge, 207
 links unbeschränkt, 341
 Linkseindeutigkeit, 56
 Liste, 29
 Liste der ZFC-Axiome, 485
 logisches Axiom, 459
 Lücke, 346

Mächtigkeit, 70, 78, 162, 265
 Mächtigkeit der reellen Funktionen, 139
 Mächtigkeit offener Mengen, 214
 mager, 380
 Mahlo-Kardinalzahl, 328
 Maß, 335
 Maß Null, 382
 mathematisches Objekt, 26
 Maximalitätsprinzip, 86
 mehrdimensionale Kontinua, 130
 Menge, 15
 Mengen und Klassen, 468
 Mengenbildung über Eigenschaften, 31
 Mengenlehre als Rahmentheorie, 43
 Mengensysteme, 40
 mengentheoretische Eigenschaften, 452
 meßbare Kardinalzahl, 331
 Metaebene, 447
 Metamathematik, 45, 447
 Minimallogik, 458
 Mirimanovsches Paradoxon, 185
 Modelle, 153
 modulo, 52
 Modus Ponens, 453
 Morse-Kelley, 440, 474
 Multiplikation von Kardinalzahlen, 164, 265
 Multiplikation linearer Ordnungen, 288
 Multiplikation zweier Ordnungstypen, 289
 Multiplikationssatz, 175, 243, 298, 301, 304

Nach oben beschränkt, 207
 nach unten beschränkt, 207
 Nachfolger einer Kardinalzahl, 171, 266, 267
 Nachfolger einer Ordinalzahl, 254, 260

Nachfolger (lineare Ordnung), 346
 Nachfolgerelement, 226
 Nachfolgerordinalzahl, 254, 260
 Nachfolgerstufe, 87
 naives Komprehensionsprinzip, 33
 natürliche Approximation von $\mathbb{C}$, 377
 natürliche Exponentiation, 302
 natürliche Zahlen, 27
 Neumann-Zermelo-Stufen, 190, 276
 nichtstationär, 322
 nirgends dicht, 380
 normaler Filter, 323
 Normalform, 299
 Normalfunktion, 312
 $\sigma\delta$ -Notation, 279
 n-reflexiv, 185
 Nullfolge, 356

Obere Schranke, 207, 245, 345
 Obermenge, 33
 Objekte, 26
 offen, 214
 Operation, 39, 271
 $\mathcal{A}$ -Operation, 397
 Operationen mit Klassen, 470
 Orbit, 97
 ordinale Wohlordnung, 257
 Ordinalzahl, 87, 252, 257
 Ordinalzahlparadoxon, 401
 $\in$ -Ordnung, 259
 Ordnung auf den Ordinalzahlen, 253
 ordnungsisomorph, 230
 Ordnungsisomorphismus, 284, 230
 ordnungstreue Abbildung, 232
 Ordnungstyp der Cantormenge, 379
 Ordnungstyp einer Wohlordnung, 262

Paarbildung, 57
 Paare, 472
 Paarmenge, 30
 Paarmengenaxiom, 423, 455
 Paarungsfunktion, 112
 Paarungssumme, 301
 Paradoxien, 183, 401
 Parameter, 32
 partielle Ordnung, 225, 244
 perfekt, 213
 perfekt reduzierbar, 392
 perfekter Kern, 360
 Platonismus, 21, 466
 potentiell unendlich, 23
 Potenzmenge, 41
 Potenzmengenaxiom, 44, 428, 455
 prädikabel, 192
 Prädikatenlogik erster Stufe, 450
 Primformel, 450

Primideal, 319
 Prinzip der Elimination, 473
 Produkt, 167, 288
 Projektion, 396
 projektiv, 397
 projektive Determiniertheit, 398
 Punktnotation, 451

Quantoraxiom, 459
 Quantorenzeichen, 449

Rang, 190
 rationale Zahlen, 27
 rationales Intervall, 361
 rechts unbeschränkt, 341
 Rechtseindeutigkeit, 54
 reductio ad absurdum, 458
 reduktibel, 365
 reduzibel, 373
 reelle Funktion, 139
 reelle Intervalle, 208
 reelle Zahlen, 28
 reflexiv, 52
 Reflexivität, 67
 Regel, 453
 regressiv, 325
 regulär, 307
 Regularitätsaxiom, 436
 Regularitätseigenschaft, 398
 Rekursion, 71, 271
 Relation, 51, 471
 relationale Klasse, 471
 Relationszeichen, 449
 relative Komplemente, 36
 Repräsentant, 52
 Repräsentationssatz, 261
 Residuum einer Punktmenge, 371
 Russell-Zermelosches Paradoxon, 183

Saturiert, 327
 Satz, 451
 Satz von Cantor, 145
 Satz von Cantor-Bernstein, 71
 Satz von Fueter-Polya, 113
 Satz von König-Zermelo, 170
 Scheeffer-Eigenschaft, 392
 Schema, 425
 Schichtung von V , 190
 Schmidt-Expansion, 85
 Schnitt, 35, 345
 Scholastik, 68
 Schuhpaare, 438
 schwach aufsteigend, 312
 schwach unerreichbare Kardinalzahl, 313
 semantische Paradoxien, 191
 separabel, 352

singulär, 307
 singuläre Kardinalzahlhypothese, 310
 Spiralaufzählung, 116
 Sprungstelle, 346
 Stabilität, 458
 starke Limeskardinalzahl, 310
 stationär, 324
 stetig (Ordinalzahlfolge), 312
 stetig (reelle Funktion), 140
 stetige Surjektion von $[0, 1]$ nach $[0, 1]^2$, 142
 strikt aufsteigend, 312
 Struktur, 224
 Substitution, 452
 Subtraktion, 35
 Subtraktion abzählbarer Mengen, 127
 Summe, 167
 Summe von Ordnungstypen, 290
 Supremum in $\mathbb{R}$, 123, 207
 Supremum, kardinales, 171
 Supremum (lineare Ordnung), 345
 Supremum (Ordinalzahlen), 255, 260
 surjektiv, 56
 Suslin-Hypothese, 357
 Suslinsche Menge, 396
 Symbolmenge, 450
 Symmetrie, 67
 symmetrisch, 52
 symmetrische Differenz, 38

Tarski-endlich, 105
 Taubenschlagprinzip, 102
 ω -te Ableitung, 217
 Teichmüller-Tukey Lemma, 248
 Teilmenge, 33
 tertium non datur, 18
 Theoremschema, 473
 Träger, 224
 transfinite Folge, 255
 transitive Menge, 258
 transitive Relation, 34, 52
 Translationsinvarianz, 336
 Transversalfunktionen, 167
 transzendent, 351
 transzendente Zahlen, 122
 Trick von Julius König, 132
 Trick von Scott, 277
 Tupel, 50

Ueberabzählbar, 120
 Überabzählbarkeit von $\mathbb{R}$, 121, 123
 Ulam-Matrix, 326
 Ultrafilter, 318
 ε -Umgebung, 208
 Umgebung, 208
 Umkehrfunktion, 58, 472
 unabhängig, 152

Unabhängigkeitsbeweise, 152
 unbeschränkt, 319, 341
 Und-Einführung, 453
 unendlich, 93
 unendliche Kardinalitäten, 163
 Unendlichkeit, 93
 Unendlichkeit und natürliche Zahlen, 97
 Unendlichkeitsaxiom, 427, 431, 455
 Unendlichkeitssymbole, 218
 unerreichbare Kardinalzahl, 314
 uniformes System, 395
 Universum, 469
 untere Schranke, 207
 Unzerlegbarkeit von $|\mathbb{R}|$, 178
 Urbild, 61
 Urelemente, 26

Variable, 32
 Variablenzeichen, 449
 Verallgemeinerte Kontinuumshypothese, 157
 Vereinigung, 35
 Vereinigung von Wohlordnungen, 227
 Vereinigungsmengenaxiom, 423, 455
 Vergleichbarkeitssatz, 81
 Verkettung, 59, 472
 Verkettung zur Identität, 60
 Verknüpfung, 59
 Vervollständigung, 354
 Vitali-Hausdorff Menge, 336
 $V = L$, 151
 vollständig ($\mathbb{R}$), 207
 vollständig (Filter, Ideal), 321
 vollständige lineare Ordnung, 345
 vollständiges Repräsentantensystem, 52
 Vollständigkeit der reellen Zahlen, 123
 von-Neumann-Bernays-Gödel, 474
 von-Neumann-Zermelo-Stufen, 190, 276
 Vorgänger, 346

Wegfall der Paradoxien, 426
 Wertebereich, 51
 Wohlordnung, 225
 Wohlordnungen gleicher Länge, 230
 Wohlordnungssatz, 81

Zahlreihe Z_0 von Zermelo, 429
 Zeichen, 449
 Zerlegung, 53, 178
 Zermelosystem, 85, 245
 Ziel in einem Zermelosystem, 85
 Ziel in einer partiellen Ordnung, 245
 Zorn-Zermelo, Satz von, 84, 86, 245
 Zuordnung, 54
 Zusammenfassung, 15
 zyklisch, 97