

Grundkurs Mathematik

LEHRBUCH

Otto Forster

Analysis 1

Differential-
und Integralrechnung
einer Veränderlichen

11. Auflage



Springer Spektrum

Grundkurs Mathematik

Berater

Prof. Dr. Martin Aigner,
Prof. Dr. Peter Gritzmann,
Prof. Dr. Volker Mehrmann,
Prof. Dr. Gisbert Wüstholtz

Die Reihe „Grundkurs Mathematik“ ist die bekannte Lehrbuchreihe im handlichen kleinen Taschenbuch-Format passend zu den mathematischen Grundvorlesungen, vorwiegend im ersten Studienjahr. Die Bücher sind didaktisch gut aufbereitet, kompakt geschrieben und enthalten viele Beispiele und Übungsaufgaben.

In der Reihe werden Lehr- und Übungsbücher veröffentlicht, die bei der Klausurvorbereitung unterstützen. Zielgruppe sind Studierende der Mathematik aller Studiengänge, Studierende der Informatik, Naturwissenschaften und Technik, sowie interessierte Schülerinnen und Schüler der Sekundarstufe II.

Die Reihe existiert seit 1975 und enthält die klassischen Bestseller von Otto Forster und Gerd Fischer zur Analysis und Linearen Algebra in aktualisierter Neuauflage.

Otto Forster

Analysis 1

Differential- und Integralrechnung
einer Veränderlichen

11., erweiterte Auflage



Springer Spektrum

Prof. Dr. Otto Forster
Ludwig-Maximilians-Universität
München, Deutschland
forster@mathematik.uni-muenchen.de

ISBN 978-3-658-00316-6

ISBN 978-3-658-00317-3 (eBook)

DOI 10.1007/978-3-658-00317-3

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

Springer Spektrum

© Springer Fachmedien Wiesbaden 1976 ... 2004, 2006, 2008, 2011, 2013

Das Werk einschließlich aller seiner Teile ist urheberrechtlich geschützt. Jede Verwertung, die nicht ausdrücklich vom Urheberrechtsgesetz zugelassen ist, bedarf der vorherigen Zustimmung des Verlags. Das gilt insbesondere für Vervielfältigungen, Bearbeitungen, Übersetzungen, Mikroverfilmungen und die Einspeicherung und Verarbeitung in elektronischen Systemen.

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.

Planung/Lektorat: Ulrike Schmickler-Hirzebruch | Barbara Gerlach

Gedruckt auf säurefreiem und chlorfrei gebleichtem Papier

Springer Spektrum ist eine Marke von Springer DE.

Springer DE ist Teil der Fachverlagsgruppe Springer Science+Business Media

www.springer-spektrum.de

Vorwort zur ersten Auflage

Dieses Buch ist entstanden aus der Ausarbeitung einer Vorlesung, die ich im WS 1970/71 für Studenten der Mathematik und Physik des ersten Semesters an der Universität Regensburg gehalten habe. Diese Ausarbeitung wurde später von verschiedenen Kollegen als Begleittext zur Vorlesung benutzt.

Der Inhalt umfaßt im wesentlichen den traditionellen Lehrstoff der Analysis-Kurse des ersten Semesters an deutschen Universitäten und Technischen Hochschulen. Bei der Stoffauswahl wurde angestrebt, dem konkreten mathematischen Inhalt, der auch für die Anwendungen wichtig ist, vor einem großen abstrakten Begriffsapparat den Vorzug zu geben und dabei gleichzeitig in systematischer Weise möglichst einfach und schnell zu den grundlegenden Begriffen (Grenzwert, Stetigkeit, Differentiation, Riemannsches Integral) vorzudringen und sie mit vielen Beispielen zu illustrieren. Deshalb wurde auch die Einführung der elementaren Funktionen vor die Abschnitte über Differentiation und Integration gezogen, um dort genügend Beispielmateriale zur Verfügung zu haben. Auf die numerische Seite der Analysis (Approximation von Größen, die nicht in endlich vielen Schritten berechnet werden können) wird an verschiedenen Stellen eingegangen, um den Grenzwertbegriff konkreter zu machen.

Der Umfang des Stoffes ist so angelegt, daß er in einer vierstündigen Vorlesung in einem Wintersemester durchgenommen werden kann. Die einzelnen Paragraphen entsprechen je nach Länge einer bis zwei Vorlesungs-Doppelstunden. Bei Zeitmangel können die §§ 17 und 23 sowie Teile der §§ 16 (Konvexität) und 20 (Gamma-Funktion) weggelassen werden.

Für seine Unterstützung möchte ich mich bei Herrn D. Leistner bedanken. Er hat die seinerzeitige Vorlesungs-Ausarbeitung geschrieben, beim Lesen der Korrekturen geholfen und das Namens- und Sachverzeichnis erstellt.

Münster, Oktober 1975

O. Forster

Vorwort zur 5. Auflage

Die erste Auflage dieses Buches erschien 1976. Seitdem hat es viele Jahrgänge von Studentinnen und Studenten der Mathematik und Physik beim Beginn ihres Analysis-Studiums begleitet. Aufgrund der damaligen Satz-Technik waren bei Neuauflagen nur geringfügige Änderungen möglich. Die einzige wesentliche Neuerung war das Erscheinen des Übungsbuchs zur Analysis 1 [FW].

Bei der jetzigen Neuauflage erhielt der Text nicht nur eine neue äußere Form ($\text{\TeX}$ -Satz), sondern wurde auch gründlich überarbeitet, um ihn wo möglich noch verständlicher zu machen. An verschiedenen Stellen wurden Bezüge zur Informatik hergestellt. So erhielt §5, in dem u.a. die Entwicklung reeller Zahlen als Dezimalbrüche (und allgemeiner b -adische Brüche) behandelt wird, einen Anhang über die Darstellung reeller Zahlen im Computer. In §9 finden sich einige grundsätzliche Bemerkungen zur Berechenbarkeit reeller Zahlen. Verschiedene numerische Beispiele wurden durch Programm-Code ergänzt, so dass die Rechnungen direkt am Computer nachvollzogen werden können. Dabei wurde der PASCAL-ähnliche Multipräzisions-Interpreter ARIBAS benutzt, den ich ursprünglich für das Buch [Fo] entwickelt habe, und der frei über das Internet erhältlich ist (Einzelheiten dazu auf Seite VIII). Die Programm-Beispiele lassen sich aber leicht auf andere Systeme, wie Maple oder Mathematica übertragen. In diesem Zusammenhang sei auch auf das Buch [BM] hingewiesen.

Insgesamt wurden aber für die Neuauflage die bewährten Charakteristiken des Buches beibehalten, nämlich ohne zu große Abstraktionen und ohne Stoffüberladung die wesentlichen Inhalte gründlich zu behandeln und sie mit konkreten Beispielen zu illustrieren. So hoffe ich, dass das Buch auch weiterhin seinen Leserinnen und Lesern den Einstieg in die Analysis erleichtern wird.

Wertvolle Hilfe habe ich von Herrn H. Stoppel erhalten. Er hat seine $\text{\TeX}$ -Erfahrung als Autor des Buches [SG] eingebracht und den Hauptteil der $\text{\TeX}$ nischen Herstellung der Neuauflage übernommen. Viele der Bilder wurden von Herrn V. Solinus erstellt. Ihnen sei herzlich gedankt, ebenso Frau Schmickler-Hirzebruch vom Vieweg-Verlag, die sich mit großem Engagement für das Zustandekommen der Neuauflage eingesetzt hat.

München, April 1999

Otto Forster

Vorwort zur 10. Auflage

Für die 10. Auflage habe ich den Text an verschiedenen Stellen weiter überarbeitet. Vor allem die Paragraphen 20 bis 23 wurden durch weitere Beispiele ergänzt (z.B. Integral-Sinus, Partialbruch-Zerlegung des Cotangens, Sinus-Produkt, Bernoulli-Zahlen und -Polynome, Fresnel-Integrale). Natürlich kann dieses Material aus Zeitmangel in der Vorlesung nur teilweise oder gar nicht gebracht werden; es eignet sich aber gut zum Selbststudium oder für Proseminare.

München, März 2011

Otto Forster

Vorwort zur 11. Auflage

Bei der Überarbeitung zur 11. Auflage ergaben sich keine größeren Änderungen; an einigen Stellen wurden neue Beispiele und Aufgaben hinzugefügt.

München, Juni 2012

Otto Forster

Software zum Buch

Die Programm-Beispiele des Buches sind für ARIBAS geschrieben. Dies ist ein Multipräzisions-Interpreter mit einer PASCAL-ähnlichen Syntax. Er ist (unter der GNU General Public License) frei über das Internet erhältlich. Es gibt Versionen von ARIBAS für verschiedene Plattformen, wie MsWindows (Windows XP, Windows 7), LINUX, Mac OS X und andere UNIX-Systeme. Für diejenigen, die hinter die Kulissen sehen wollen, ist auch der C-Source-Code von ARIBAS verfügbar. Um ARIBAS zu erhalten, gehe man auf die WWW-Homepage des Verfassers,

`http://www.mathematik.uni-muenchen.de/~forster`

und von dort zum Unterpunkt Software/ARIBAS. Dort finden sich weitere Informationen. Da ARIBAS ein kompaktes System ist, muss nur etwa 1/4 MB heruntergeladen werden.

Von der oben genannten Homepage gelangt man über den Unterpunkt Bücher/Analysis auch zur Homepage dieses Buches. Von dort sind die Listings der Programm-Beispiele erhältlich, so dass sie nicht mühsam abgetippt werden müssen. Im Laufe der Zeit werden noch weitere Listings zu numerischen Übungsaufgaben und zu Ergänzungen zum Text dazukommen. Ebenfalls wird dort eine Liste der unvermeidlich zutage tretenden *Errata* abgelegt werden. Die aufmerksamen Leserinnen und Leser seien ermuntert, mir Fehler per E-Mail an folgende Adresse zu melden:

`forster@mathematik.uni-muenchen.de`

Inhaltsverzeichnis

1	Vollständige Induktion	1
2	Die Körper-Axiome	17
3	Die Anordnungs-Axiome	25
4	Folgen, Grenzwerte	35
5	Das Vollständigkeits-Axiom	49
6	Wurzeln	62
7	Konvergenz-Kriterien für Reihen	70
8	Die Exponentialreihe	83
9	Punktmengen	90
10	Funktionen. Stetigkeit	103
11	Sätze über stetige Funktionen	113
12	Logarithmus und allgemeine Potenz	124
13	Die Exponentialfunktion im Komplexen	136
14	Trigonometrische Funktionen	145
15	Differentiation	163
16	Lokale Extrema. Mittelwertsatz. Konvexität	177
17	Numerische Lösung von Gleichungen	192
18	Das Riemannsche Integral	202
19	Integration und Differentiation	217
20	Uneigentliche Integrale. Die Gamma-Funktion	235
21	Gleichmäßige Konvergenz von Funktionenfolgen	255
22	Taylor-Reihen	279
23	Fourier-Reihen	304
	Zusammenstellung der Axiome der reellen Zahlen	323
	Literaturhinweise	324
	Namens- und Sachverzeichnis	326
	Symbolverzeichnis	332

§ 1 Vollständige Induktion

Der Beweis durch vollständige Induktion ist ein wichtiges Hilfsmittel in der Mathematik. Es kann häufig bei Problemen folgender Art angewandt werden: Es soll eine Aussage $A(n)$ bewiesen werden, die von einer natürlichen Zahl $n \geq 1$ abhängt. Dies sind in Wirklichkeit unendlich viele Aussagen $A(1), A(2), A(3), \dots$, die nicht alle einzeln bewiesen werden können. Hier hilft vollständige Induktion, die unter geeigneten Umständen erlaubt, in endlich vielen Schritten unendlich viele Aussagen zu beweisen.

Beweisprinzip der vollständigen Induktion

Sei n_0 eine ganze Zahl und $A(n)$ für jede ganze Zahl $n \geq n_0$ eine Aussage. Um $A(n)$ für alle $n \geq n_0$ zu beweisen, genügt es, zu zeigen:

(I.0) $A(n_0)$ ist richtig (Induktions-Anfang).

(I.1) Für ein beliebiges $n \geq n_0$ gilt:

Falls $A(n)$ richtig ist, so ist auch $A(n+1)$ richtig (Induktions-Schritt).

Die Wirkungsweise dieses Beweisprinzips ist leicht einzusehen: Nach (I.0) ist zunächst $A(n_0)$ richtig. Wendet man (I.1) auf den Fall $n = n_0$ an, erhält man die Gültigkeit von $A(n_0 + 1)$. Wiederholte Anwendung von (I.1) liefert dann die Richtigkeit von $A(n_0 + 2), A(n_0 + 3), \dots$, usw.

Als erstes Beispiel beweisen wir damit eine nützliche Formel für die Summe der ersten n natürlichen Zahlen.

Satz 1. Für jede natürliche Zahl n gilt:

$$1 + 2 + 3 + \dots + n = \frac{n(n+1)}{2}.$$

Beweis. Wir setzen zur Abkürzung $S(n) = 1 + 2 + \dots + n$ und zeigen die Gleichung $S(n) = \frac{n(n+1)}{2}$ durch vollständige Induktion.

Induktions-Anfang $n = 1$. Es ist $S(1) = 1$ und $\frac{1(1+1)}{2} = 1$, also gilt die Formel für $n = 1$.

Induktions-Schritt $n \rightarrow n+1$. Wir nehmen an, dass $S(n) = \frac{n(n+1)}{2}$ gilt (Induktions-Voraussetzung) und müssen zeigen, dass daraus die Formel $S(n+1) = \frac{(n+1)(n+2)}{2}$

folgt. Dies sieht man so:

$$\begin{aligned} S(n+1) &= S(n) + (n+1) \stackrel{\text{IV}}{=} \frac{n(n+1)}{2} + n+1 \\ &= \frac{(n+1)(n+2)}{2}, \quad \text{q.e.d.} \end{aligned}$$

Dabei deutet $\stackrel{\text{IV}}{=}$ an, dass an dieser Stelle die Induktions-Voraussetzung benutzt wurde.

Der Satz 1 erinnert an die bekannte Geschichte über C.F. Gauß, der als kleiner Schüler seinen Lehrer dadurch in Erstaunen versetzte, dass er die Aufgabe, die Zahlen von 1 bis 100 zusammenzuzählen, in kürzester Zeit im Kopf löste. Gauß verwendete dazu keine vollständige Induktion, sondern benutzte folgenden Trick: Er fasste den ersten mit dem letzten Summanden, den zweiten mit dem vorletzten zusammen, usw.

$$\begin{aligned} 1 + 2 + \dots + 100 &= (1 + 100) + (2 + 99) + \dots + (50 + 51) \\ &= 50 \cdot 101 = 5050. \end{aligned}$$

Natürlich ergibt sich dasselbe Resultat mit der Formel aus Satz 1.

Summenzeichen. Formeln wie in Satz 1 lassen sich oft prägnanter unter Verwendung des Summenzeichens schreiben. Seien $m \leq n$ ganze Zahlen. Für jede ganze Zahl k mit $m \leq k \leq n$ sei a_k eine reelle Zahl. Dann setzt man

$$\sum_{k=m}^n a_k := a_m + a_{m+1} + \dots + a_n.$$

(Dabei bedeutet $X := A$, dass X nach Definition gleich A ist.) Für $m = n$ besteht die Summe aus dem einzigen Summanden a_m . Es ist zweckmäßig, für $n = m - 1$ folgende Konvention einzuführen:

$$\sum_{k=m}^{m-1} a_k := 0 \quad (\text{leere Summe}).$$

Man kann die etwas unbefriedigenden Pünktchen $\dots$ in der Definition des Summenzeichens vermeiden, wenn man *Definition durch vollständige Induktion* benutzt: Für den Induktions-Anfang setzt man $\sum_{k=m}^{m-1} a_k := 0$ und verwendet als Induktionsschritt

$$\sum_{k=m}^{n+1} a_k := \left(\sum_{k=m}^n a_k \right) + a_{n+1} \quad \text{für alle } n \geq m-1.$$

Als natürliche Zahlen bezeichnen wir alle Elemente der Menge

$$\mathbb{N} := \{0, 1, 2, 3, \dots\}$$

der nicht-negativen ganzen Zahlen (einschließlich der Null). Mit

$$\mathbb{Z} := \{0, \pm 1, \pm 2, \pm 3, \dots\}$$

wird die Menge aller ganzen Zahlen bezeichnet.

Nun lässt sich Satz 1 so aussprechen: Es gilt

$$\sum_{k=1}^n k = \frac{n(n+1)}{2} \quad \text{für alle } n \in \mathbb{N}.$$

(Für $n = 0$ gilt die Formel trivialerweise, da beide Seiten der Gleichung gleich null sind.)

(1.1) Als weiteres Beispiel für einen Beweis durch vollständige Induktion betrachten wir die Summe der ersten ungeraden natürlichen Zahlen:

$$\begin{aligned} 1 &= 1, \\ 1 + 3 &= 4, \\ 1 + 3 + 5 &= 9, \\ 1 + 3 + 5 + 7 &= 16, \\ 1 + 3 + 5 + 7 + 9 &= 25, \\ &\dots \end{aligned}$$

Es fällt auf, dass sich stets eine Quadratzahl ergibt. Um zu zeigen, dass dies allgemein richtig ist, verwenden wir wieder vollständige Induktion.

Satz 2. Für alle natürlichen Zahlen n gilt $\sum_{k=1}^n (2k-1) = n^2$.

Beweis. Induktions-Anfang $n = 0$.

$$\sum_{k=1}^0 (2k-1) = 0 = 0^2 \quad (\text{leere Summe}).$$

Wem das Hantieren mit leeren Summen nicht geheuer ist, der kann auch mit $n = 1$ anfangen:

$$\sum_{k=1}^1 (2k-1) = 2 \cdot 1 - 1 = 1 = 1^2.$$

Induktions-Schritt $n \rightarrow n + 1$.

$$\begin{aligned}\sum_{k=1}^{n+1} (2k-1) &= \sum_{k=1}^n (2k-1) + (2(n+1)-1) \stackrel{\text{IV}}{=} n^2 + 2n + 1 \\ &= (n+1)^2, \quad \text{q.e.d.}\end{aligned}$$

(1.2) Primzahlformel? Dass die Verhältnisse nicht immer so einfach liegen, wie in Beispiel (1.1), sollen die folgenden Betrachtungen zeigen. Das Polynom $P(x)$ sei wie folgt definiert:

$$P(x) := x^2 + x + 41.$$

Wir berechnen die Werte dieses Polynoms an den natürlichen Zahlen: Für $0 \leq n \leq 10$ erhält man

n	0	1	2	3	4	5	6	7	8	9	10
$P(n)$	41	43	47	53	61	71	83	97	113	131	151

Es fällt auf, dass lauter Primzahlen¹ entstehen. (Um sich z.B. davon zu überzeugen, dass 151 prim ist, muss man, da $13^2 > 151$, nur prüfen, dass 151 durch keine der Primzahlen 2, 3, 5, 7, 11 teilbar ist.) Man könnte deshalb vermuten, dass $P(n)$ für alle natürlichen Zahlen n prim ist. Ein Beweis bietet sich jedoch nicht unmittelbar an. Aus Satz 1 folgt

$$P(n) = 41 + 2 \sum_{k=1}^n k.$$

Da die Summe einer ungeraden und einer geraden Zahl ungerade ist, ist also $P(n)$ stets ungerade. Dies ist dafür, dass eine ganze Zahl > 2 prim ist, zwar eine notwendige, aber keineswegs hinreichende Bedingung. Etwas weitergehend wollen wir zeigen, dass keine der Zahlen $P(n)$ durch eine der Primzahlen $p = 2, 3, 5, 7, 11$ teilbar ist. Hierfür geben wir einen Beweis durch vollständige Induktion in einer modifizierten Form. Es bezeichne $A(n)$ die folgende Aussage: " $P(n)$ ist durch keine der Primzahlen 2, 3, 5, 7, 11 teilbar." Dann sind $A(0), A(1), \dots, A(10)$ wahr, da die $P(n)$, $n \leq 10$, Primzahlen > 11 sind. Dies ist unser Induktions-Anfang. Sei nun $n > 10$. Wir nehmen an, dass $A(k)$ für

¹Bekanntlich heißt eine ganze Zahl $N > 1$ Primzahl, wenn sie keinen (ganzzahligen) Teiler t mit $1 < t < N$ besitzt, d.h. wenn keine ganzen Zahlen $s, t > 1$ existieren mit $N = t \cdot s$. Die Zahl 1 ist definitionsgemäß keine Primzahl. Da für ein zusammengesetztes $N = t \cdot s$ für den kleineren Faktor (o.B.d.A. sei $t \leq s$) gilt $t^2 \leq N$, folgt: Eine ganze Zahl $N > 1$ ist genau dann prim, wenn sie keinen Primfaktor $p \leq \sqrt{N}$ besitzt.

alle $k < n$ wahr ist (Induktions-Voraussetzung) und wollen daraus ableiten (Induktions-Schritt), dass dann auch $A(n)$ gilt.²

Für die Gültigkeit von $A(n)$ ist zu zeigen, dass $P(n)$ durch keine der Primzahlen $p \in \{2, 3, 5, 7, 11\}$ teilbar ist. Dazu berechnen wir die Differenz

$$\begin{aligned} P(n) - P(n-p) &= n^2 + n - (n-p)^2 - (n-p) \\ &= 2np - p^2 + p = (2n - p + 1)p. \end{aligned}$$

Da $n \geq 11$ und $p \leq 11$, gilt $0 \leq n - p < n$. Nach Induktions-Voraussetzung ist deshalb $A(n-p)$ wahr, also $P(n-p)$ nicht durch p teilbar. Daher lässt es sich schreiben als $P(n-p) = mp + r$ mit ganzen Zahlen m, r und $0 < r < p$. Daraus folgt

$$P(n) = (m + 2n - p + 1)p + r,$$

was zeigt, dass $P(n)$ nicht durch p teilbar ist, q.e.d.

Damit haben wir bewiesen, dass die unendlich vielen Zahlen $P(n)$, $n \in \mathbb{N}$, unter ihnen z.B.

$$P(100) = 10141 \quad \text{und} \quad P(1000) = 1001041,$$

alle keinen Primfaktor $p \leq 11$ besitzen. Dies beweist natürlich noch nicht, dass diese Zahlen prim sind. Tatsächlich sind $P(100)$ und $P(1000)$ jedoch (zufällig?) Primzahlen. Im Fall von $P(1000)$ hat man dazu nachzuprüfen (am besten mit Computer-Hilfe), dass diese Zahl keinen Primfaktor ≤ 1000 hat. Wir erweitern unsere Tabelle für $P(n)$ noch um einige Werte:

n	0	1	2	3	4	5	6	7	8	9
$P(n)$	41	43	47	53	61	71	83	97	113	131
$P(10+n)$	151	173	197	223	251	281	313	347	383	421
$P(20+n)$	461	503	547	593	641	691	743	797	853	911
$P(30+n)$	971	1033	1097	1163	1231	1301	1373	1447	1523	1601

Wieder entstehen lauter Primzahlen. Jedoch

$$P(40) = 40(40+1) + 41 = 41^2$$

ist keine Primzahl. Die Vermutung "Für alle natürlichen Zahlen n ist $P(n) = n^2 + n + 41$ eine Primzahl" ist daher falsch, obwohl sie in den Spezialfällen

²Dass das hier benutzte Induktions-Schema mit dem eingangs dieses Paragraphen beschriebenen Beweisprinzip äquivalent ist, sieht man so: Wir bezeichnen mit $A^*(n)$ die Aussage: "A(k) gilt für alle natürlichen Zahlen $k \leq n$." Dann ist $A^*(10)$ der Induktions-Anfang und der Induktions-Schritt ist die Implikation: Aus $A^*(n-1)$ folgt $A^*(n)$.

$n = 0, 1, 2, 3, 4, 5, \dots, 37, 38, 39$ zutrifft. In der Mathematik ist ein Satz falsch, wenn es auch nur ein einziges Gegenbeispiel dazu gibt. Das Sprichwort “Die Ausnahme bestätigt die Regel” ist hier nicht anwendbar. In unserem speziellen Fall gibt es sogar unendlich viele Ausnahmen, denn es ist leicht zu sehen, dass für alle n der Gestalt $n = 41m$ die Zahl $P(n)$ durch 41 teilbar ist (darauf hätte man natürlich gleich kommen können). Es lässt sich sogar zeigen, dass es überhaupt kein Polynom $P(x)$ vom Grad ≥ 1 mit ganzzahligen Koeffizienten gibt, so dass $P(n)$ für alle natürlichen Zahlen n eine Primzahl ist.

Man wird sich aber fragen, ob es einen tieferen Grund dafür gibt, dass das quadratische Polynom $n^2 + n + 41$ für so viele konsekutive Werte von n eine Primzahl liefert (diese Tatsache wurde bereits von Euler entdeckt). Für eine Antwort, die weitergehende zahlentheoretische Hilfsmittel erfordert, verweisen wir die interessierte Leserin auf das Buch von Ribenboim [Ri], Kap. 5.

Wir kommen jetzt zur Definition der Fakultät einer natürlichen Zahl, die in der Kombinatorik eine wichtige Rolle spielt.

Definition (Fakultät). Für $n \in \mathbb{N}$ setzt man

$$n! := \prod_{k=1}^n k = 1 \cdot 2 \cdot \dots \cdot n.$$

Das Produktzeichen ist ganz analog zum Summenzeichen definiert. Man setzt (Induktions-Anfang)

$$\prod_{k=m}^{m-1} a_k := 1 \quad (\text{leeres Produkt}),$$

und (Induktions-Schritt)

$$\prod_{k=m}^{n+1} a_k := \left(\prod_{k=m}^n a_k \right) a_{n+1} \quad \text{für alle } n \geq m-1.$$

(Das leere Produkt wird deshalb als 1 definiert, da die Multiplikation mit 1 dieselbe Wirkung hat wie wenn man überhaupt nicht multipliziert.)

Insbesondere ist $0! = 1$, $1! = 1$, $2! = 2$, $3! = 6$, $4! = 24$, $5! = 120$, ...

Satz 3. Die Anzahl aller möglichen Anordnungen einer n -elementigen Menge $\{A_1, A_2, \dots, A_n\}$ ist gleich $n!$.

Beweis durch vollständige Induktion.

Induktions-Anfang $n = 1$. Eine einelementige Menge besitzt nur eine Anordnung ihrer Elemente. Andererseits ist $1!$ ebenfalls gleich 1.

Induktions-Schritt $n \rightarrow n + 1$. Die möglichen Anordnungen der $(n + 1)$ -elementigen Menge $\{A_1, A_2, \dots, A_{n+1}\}$ zerfallen folgendermaßen in $n + 1$ Klassen C_k , $k = 1, \dots, n + 1$: Die Anordnungen der Klasse C_k haben das Element A_k an erster Stelle, bei beliebiger Anordnung der übrigen n Elemente. Nach Induktions-Voraussetzung besteht jede Klasse aus $n!$ Anordnungen. Die Gesamtzahl aller möglichen Anordnungen von $\{A_1, A_2, \dots, A_{n+1}\}$ ist also gleich $(n + 1)n! = (n + 1)!$, q.e.d.

Bemerkung. Die beim Induktions-Schritt benützte Überlegung kann man dazu verwenden, alle Anordnungen systematisch aufzuzählen (wir schreiben kurz k statt A_k).

$n = 2$

1 2 2 1

$n = 3$

1 2 3 2 1 3 3 1 2
1 3 2 2 3 1 3 2 1

$n = 4$

1 2 3 4	2 1 3 4	3 1 2 4	4 1 2 3
1 2 4 3	2 1 4 3	3 1 4 2	4 1 3 2
1 3 2 4	2 3 1 4	3 2 1 4	4 2 1 3
1 3 4 2	2 3 4 1	3 2 4 1	4 2 3 1
1 4 2 3	2 4 1 3	3 4 1 2	4 3 1 2
1 4 3 2	2 4 3 1	3 4 2 1	4 3 2 1

Definition. Für natürliche Zahlen n und k setzt man

$$\binom{n}{k} = \prod_{j=1}^k \frac{n-j+1}{j} = \frac{n(n-1) \cdot \dots \cdot (n-k+1)}{1 \cdot 2 \cdot \dots \cdot k}.$$

Die Zahlen $\binom{n}{k}$ heißen *Binomial-Koeffizienten* wegen ihres Auftretens im binomischen Lehrsatz (vgl. den folgenden Satz 5).

Aus der Definition folgt unmittelbar

$$\binom{n}{0} = 1, \quad \binom{n}{1} = n \quad \text{für alle } n \geq 0,$$

$$\binom{n}{k} = 0 \quad \text{für } k > n, \quad \text{sowie}$$

$$\binom{n}{k} = \frac{n!}{k!(n-k)!} = \binom{n}{n-k} \quad \text{für } 0 \leq k \leq n.$$

Definiert man noch $\binom{n}{k} := 0$ für $k < 0$, so gilt

$$\binom{n}{k} = \binom{n}{n-k} \quad \text{für alle } n \in \mathbb{N} \text{ und } k \in \mathbb{Z}.$$

Hilfssatz. Für alle natürlichen Zahlen $n \geq 1$ und alle $k \in \mathbb{Z}$ gilt

$$\binom{n}{k} = \binom{n-1}{k-1} + \binom{n-1}{k}.$$

Beweis. Für $k \geq n$ und $k \leq 0$ verifiziert man die Formel unmittelbar. Es bleibt also der Fall $0 < k < n$ zu betrachten. Dann ist

$$\begin{aligned} \binom{n-1}{k-1} + \binom{n-1}{k} &= \frac{(n-1)!}{(k-1)!(n-k)!} + \frac{(n-1)!}{k!(n-k-1)!} \\ &= \frac{k(n-1)! + (n-k)(n-1)!}{k!(n-k)!} = \frac{n(n-1)!}{k!(n-k)!} = \binom{n}{k}. \end{aligned}$$

Satz 4. Die Anzahl der k -elementigen Teilmengen einer n -elementigen Menge $\{A_1, A_2, \dots, A_n\}$ ist gleich $\binom{n}{k}$.

Bemerkung. Daraus folgt auch, dass die Zahlen $\binom{n}{k}$ ganz sind, was aus ihrer Definition nicht unmittelbar ersichtlich ist.

Beweis. Wir beweisen die Behauptung durch vollständige Induktion nach n .

Induktions-Anfang $n = 1$. Die Menge $\{A_1\}$ besitzt genau eine nullelementige Teilmenge, nämlich die leere Menge $\emptyset$, und genau eine einelementige Teilmenge, nämlich $\{A_1\}$. Andererseits ist auch $\binom{1}{0} = \binom{1}{1} = 1$.

(Übrigens gilt der Satz auch für $n = 0$, denn die leere Menge besitzt genau eine nullelementige Teilmenge, nämlich die leere Menge, und es gilt $\binom{0}{0} = \frac{0!}{0!0!} = 1$.)

Induktions-Schritt $n \rightarrow n+1$. Die Behauptung sei für Teilmengen der n -elementigen Menge $M_n := \{A_1, \dots, A_n\}$ schon bewiesen. Wir betrachten nun die k -elementigen Teilmengen von $M_{n+1} := \{A_1, \dots, A_n, A_{n+1}\}$. Für $k = 0$ und $k = n+1$ ist die Behauptung trivial, wir dürfen also $1 \leq k \leq n$ annehmen. Jede k -elementige Teilmenge von M_{n+1} gehört zu genau einer der folgenden Klassen: $\mathcal{T}_0$ besteht aus allen k -elementigen Teilmengen von M_{n+1} , die A_{n+1} nicht

enthalten, und $\mathcal{T}_1$ aus denjenigen k -elementigen Teilmengen, die A_{n+1} enthalten. Die Anzahl der Elemente von $\mathcal{T}_0$ ist gleich der Anzahl der k -elementigen Teilmengen von M_n , also nach Induktions-Voraussetzung gleich $\binom{n}{k}$. Da die Teilmengen der Klasse $\mathcal{T}_1$ alle das Element A_{n+1} enthalten, und die übrigen $k-1$ Elemente der Menge M_n entnommen sind, besteht $\mathcal{T}_1$ nach Induktions-Voraussetzung aus $\binom{n}{k-1}$ Elementen. Insgesamt gibt es also (unter Benutzung des Hilfssatzes)

$$\binom{n}{k} + \binom{n}{k-1} = \binom{n+1}{k}$$

k -elementige Teilmengen von M_{n+1} , q.e.d.

(1.3) Beispiel. Es gibt

$$\binom{49}{6} = \frac{49 \cdot 48 \cdot 47 \cdot 46 \cdot 45 \cdot 44}{1 \cdot 2 \cdot 3 \cdot 4 \cdot 5 \cdot 6} = 13983816$$

6-elementige Teilmengen einer Menge von 49 Elementen. Die Chance, beim Lotto "6 aus 49" die richtige Kombination zu erraten, ist also etwa 1 : 14 Millionen.

Der folgende Satz verallgemeinert die schon aus der Schule bekannte Formel $(x+y)^2 = x^2 + 2xy + y^2$.

Satz 5 (Binomischer Lehrsatz). *Seien x, y reelle Zahlen und n eine natürliche Zahl. Dann gilt*

$$(x+y)^n = \sum_{k=0}^n \binom{n}{k} x^{n-k} y^k.$$

Beweis durch vollständige Induktion nach n .

Induktions-Anfang $n = 0$. Da nach Definition $a^0 = 1$ für jede reelle Zahl a (leeres Produkt), ist $(x+y)^0 = 1$ und

$$\sum_{k=0}^0 \binom{0}{k} x^{n-k} y^k = \binom{0}{0} x^0 y^0 = 1.$$

Induktions-Schritt $n \rightarrow n+1$.

$$(x+y)^{n+1} = (x+y)^n x + (x+y)^n y.$$

Für den ersten Summanden der rechten Seite erhält man unter Benutzung der

Induktions-Voraussetzung

$$(x+y)^n x = \sum_{k=0}^n \binom{n}{k} x^{n+1-k} y^k = \sum_{k=0}^{n+1} \binom{n}{k} x^{n+1-k} y^k.$$

Dabei haben wir verwendet, dass $\binom{n}{n+1} = 0$. Für die Umformung des zweiten Summanden verwenden wir die offensichtliche Regel

$$\sum_{k=0}^n a_{k+1} = \sum_{k=1}^{n+1} a_k$$

über die Indexverschiebung bei Summen.

$$(x+y)^n y = \sum_{k=0}^n \binom{n}{k} x^{n-k} y^{k+1} = \sum_{k=1}^{n+1} \binom{n}{k-1} x^{n+1-k} y^k.$$

Addiert man den Summanden $\binom{n}{-1} x^{n+1} y^0 = 0$, erhält man

$$(x+y)^n y = \sum_{k=0}^{n+1} \binom{n}{k-1} x^{n+1-k} y^k.$$

Insgesamt ergibt sich, wenn man noch $\binom{n}{k} + \binom{n}{k-1} = \binom{n+1}{k}$ benutzt (Hilfssatz),

$$\begin{aligned} (x+y)^{n+1} &= \sum_{k=0}^{n+1} \binom{n}{k} x^{n+1-k} y^k + \sum_{k=0}^{n+1} \binom{n}{k-1} x^{n+1-k} y^k \\ &= \sum_{k=0}^{n+1} \binom{n+1}{k} x^{n+1-k} y^k, \quad \text{q.e.d.} \end{aligned}$$

Für die ersten n lautet der Binomische Lehrsatz ausgeschrieben

$$(x+y)^0 = 1,$$

$$(x+y)^1 = x+y,$$

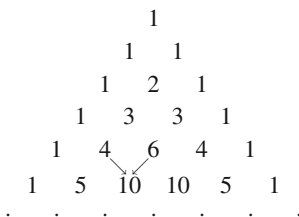
$$(x+y)^2 = x^2 + 2xy + y^2,$$

$$(x+y)^3 = x^3 + 3x^2y + 3xy^2 + y^3,$$

$$(x+y)^4 = x^4 + 4x^3y + 6x^2y^2 + 4xy^3 + y^4,$$

$$(x+y)^5 = x^5 + 5x^4y + 10x^3y^2 + 10x^2y^3 + 5xy^4 + y^5, \text{ usw.}$$

Die auftretenden Koeffizienten kann man im sog. *Pascalschen Dreieck* anordnen.



Aufgrund der Beziehung $\binom{n}{k} = \binom{n-1}{k-1} + \binom{n-1}{k}$ ist jede Zahl im Inneren des Dreiecks die Summe der beiden unmittelbar über ihr stehenden.

(1.4) Folgerungen aus dem binomischen Lehrsatz. Für alle $n \geq 1$ gilt

$$\sum_{k=0}^n \binom{n}{k} = 2^n \quad \text{und} \quad \sum_{k=0}^n (-1)^k \binom{n}{k} = 0.$$

Man erhält dies, wenn man $x = y = 1$ bzw. $x = 1, y = -1$ setzt.

Die erste dieser Formeln lässt sich nach Satz 4 kombinatorisch wie folgt interpretieren: Da $\binom{n}{k}$ die Anzahl der k -elementigen Teilmengen einer n -elementigen Menge angibt, besitzt eine n -elementige Menge insgesamt 2^n Teilmengen. Diese letztere Aussage lässt sich auch direkt beweisen: Wir ordnen jeder Teilmenge $T \subset M_n$ der n -elementigen Menge

$$M_n := \{A_1, A_2, \dots, A_n\}$$

wie folgt einen "Bit-Vektor" $(b_1, b_2, \dots, b_n)$ mit $b_i \in \{0, 1\}$ zu: Man setzt

$$b_i := \begin{cases} 1, & \text{falls } A_i \in T, \\ 0, & \text{falls } A_i \notin T. \end{cases}$$

Auf diese Weise erhält man eine umkehrbar eindeutige Zuordnung zwischen der Menge aller Teilmengen von M_n und der Menge aller n -dimensionalen Bit-Vektoren $(b_1, b_2, \dots, b_n)$. Da es für jede Komponente b_i genau zwei Möglichkeiten 0 oder 1 gibt, gibt es insgesamt 2^n solcher Vektoren, also auch 2^n Teilmengen einer n -elementigen Menge. Dies liefert nun umgekehrt einen weiteren Beweis der Formel $\sum_{k=0}^n \binom{n}{k} = 2^n$.

Satz 6 (geometrische Reihe). Für jede reelle Zahl $x \neq 1$ und jede natürliche Zahl n gilt

$$\sum_{k=0}^n x^k = \frac{1 - x^{n+1}}{1 - x}.$$

Beweis durch vollständige Induktion nach n .

Induktions-Anfang $n = 0$.

$$\sum_{k=0}^0 x^k = 1 = \frac{1 - x^{0+1}}{1 - x}.$$

Induktions-Schritt $n \rightarrow n + 1$.

$$\sum_{k=0}^{n+1} x^k = \sum_{k=0}^n x^k + x^{n+1} = \frac{1 - x^{n+1}}{1 - x} + x^{n+1} = \frac{1 - x^{(n+1)+1}}{1 - x}, \quad \text{q.e.d.}$$

Folgerungen

(1.5) Setzt man speziell $x = 2$ in Satz 6, erhält man

$$\sum_{k=0}^n 2^k = 2^{n+1} - 1,$$

insbesondere

$$\sum_{k=0}^{63} 2^k = 2^{64} - 1 = 18446\,74407\,37095\,51615.$$

Mit den Zahlen von 0 bis 63 lassen sich die Felder eines Schachbretts durchnummerieren. Der Schachfreund erinnert sich jetzt an die Legende von der Belohnung des Erfinders des Schachspiels.

(1.6) Schreibt man das Ergebnis von Satz 6 als $\sum_{k=0}^{n-1} x^k = \frac{1-x^n}{1-x}$ und multipliziert die Formel mit $x - 1$, erhält man die Zerlegung

$$x^n - 1 = (x - 1)(x^{n-1} + x^{n-2} + \dots + x + 1).$$

(Dies gilt für alle reellen Zahlen x , auch für $x = 1$, da dann beide Seiten der Gleichung den Wert 0 haben.) Daraus kann man z.B. folgendes schließen:

Eine Zahl der Gestalt $M = 2^n - 1$, (n ganze Zahl > 1), ist höchstens dann eine Primzahl, wenn n eine Primzahl ist.

Beweis hierfür. Wäre n nicht prim, könnte man n zerlegen als $n = k \cdot \ell$ mit ganzen Zahlen $k, \ell > 1$. Damit folgt

$$M = 2^{k\ell} - 1 = (2^k)^\ell - 1 = (2^k - 1)(2^{k(\ell-1)} + 2^{k(\ell-2)} + \dots + 1),$$

d.h. M ist nicht prim. Eine notwendige Bedingung dafür, dass $M = 2^n - 1$ prim ist, ist also, dass n selbst eine Primzahl ist.

Bemerkung. Solche Primzahlen wurden schon im 17. Jahrhundert von M. Mersenne studiert. Man nennt eine Primzahl der Gestalt $M_p := 2^p - 1$, (p prim), *Mersennesche Primzahl*. Die kleinsten Mersenneschen Primzahlen sind

$$\begin{aligned} M_2 &= 2^2 - 1 = 3, & M_3 &= 2^3 - 1 = 7, & M_5 &= 2^5 - 1 = 31, \\ M_7 &= 2^7 - 1 = 127, & M_{13} &= 2^{13} - 1 = 8191, & M_{17} &= 131071. \end{aligned}$$

Dagegen ist z.B. $M_{11} = 2^{11} - 1 = 2047 = 23 \cdot 89$ nicht prim. Bis heute (Juni 2012) sind erst 47 Mersennesche Primzahlen bekannt, die größte bekannte (M_p mit $p = 43112609$) hat über zwölf Millionen Dezimalstellen.

(1.7) Ersetzt man bei ungeradem $n = 2k + 1$ in der Formel von (1.6) die Variable x durch $-x$, so erhält man

$$x^{2k+1} + 1 = (x + 1)(x^{2k} - x^{2k-1} + \dots - x + 1)$$

Daraus kann man ableiten:

Eine Zahl der Gestalt $F = 2^n + 1$, (n ganze Zahl ≥ 1), ist höchstens dann eine Primzahl, wenn n eine Zweierpotenz (d.h. $n = 2^k$) ist.

Denn ist n keine Zweierpotenz, so hat n einen ungeraden Faktor, d.h. $n = \ell m$ mit einer ungeraden Zahl $\ell \geq 3$ und $m \geq 1$. Daraus folgt

$$F = (2^m)^\ell + 1 = (2^m + 1)(2^{m(\ell-1)} - 2^{m(\ell-2)} + \dots - 2^m + 1),$$

also ist F dann keine Primzahl.

Bemerkung. Die Zahlen der Gestalt $F_k := 2^{2^k} + 1$ heißen *Fermat-Zahlen*. Fermat glaubte, dass alle F_k prim seien. Dies ist richtig für

$$\begin{aligned} F_0 &= 2^1 + 1 = 3, & F_1 &= 2^2 + 1 = 5, & F_2 &= 2^4 + 1 = 17, \\ F_3 &= 2^8 + 1 = 257, & F_4 &= 2^{16} + 1 = 65537, \end{aligned}$$

aber bereits $F_5 = 2^{32} + 1 = 4294967297 = 641 \cdot 6700417$ ist eine zusammengesetzte Zahl. Bis heute wurden keine weiteren Fermatschen Primzahlen gefunden.

AUFGABEN

1.1. Lässt sich der Trick von Gauß über die Summe der ersten n natürlichen Zahlen auch anwenden, wenn n ungerade ist?

1.2. Man beweise: Für jede natürliche Zahl n hat die Zahl $P(n) := n^2 + n + 41$ keinen Primfaktor ≤ 37 .

1.3. Man beweise die Summenformeln

$$\sum_{k=1}^n k^2 = \frac{n(n+1)(2n+1)}{6} \quad \text{und} \quad \sum_{k=1}^n k^3 = \frac{n^2(n+1)^2}{4}.$$

1.4. Man finde eine Formel für

$$\sum_{k=1}^n (2k-1)^2$$

und beweise sie.

1.5. Sei r eine natürliche Zahl. Man zeige:

Es gibt rationale Zahlen $a_{r1}, \dots, a_{rr}$, so dass für alle natürlichen Zahlen n gilt

$$\sum_{k=1}^n k^r = \frac{1}{r+1} n^{r+1} + a_{rr} n^r + \dots + a_{r1} n.$$

1.6. Man beweise: Für alle natürlichen Zahlen n gilt

$$\sum_{k=1}^n \frac{1}{k(k+1)} = 1 - \frac{1}{n+1}.$$

1.7. Man beweise: Für alle natürlichen Zahlen N gilt

$$\sum_{n=1}^{2N} \frac{(-1)^{n-1}}{n} = \sum_{n=1}^N \frac{1}{N+n}.$$

1.8. Seien $a_0, a_1, \dots, a_n$ und $b_0, b_1, \dots, b_n$ reelle Zahlen und

$$A_k := \sum_{i=0}^k a_i, \quad B_k := \sum_{i=0}^k b_i \quad \text{für } k = 0, 1, \dots, n.$$

Man beweise (Abelsche partielle Summation):

$$\sum_{k=0}^n A_k b_k = A_n B_n - \sum_{k=0}^{n-1} a_{k+1} B_k.$$

1.9. Seien n, k natürliche Zahlen mit $n \geq k$. Man beweise

$$\binom{n+1}{k+1} = \sum_{m=k}^n \binom{m}{k}.$$

1.10. Man beweise: Eine n -elementige Menge ($n > 0$) besitzt ebenso viele Teilmengen mit einer geraden Zahl von Elementen wie Teilmengen mit einer ungeraden Zahl von Elementen.

1.11. Für eine reelle Zahl x und eine natürliche Zahl k werde definiert

$$\binom{x}{k} := \prod_{j=1}^k \frac{x-j+1}{j} = \frac{x(x-1) \cdot \dots \cdot (x-k+1)}{k!},$$

insbesondere $\binom{x}{0} = 1$. Man beweise für alle reellen Zahlen x, y und alle natürlichen Zahlen n

$$\binom{x+y}{n} = \sum_{k=0}^n \binom{x}{n-k} \binom{y}{k}.$$

1.12. Ersetzt man im Pascalschen Dreieck die Einträge durch kleine rechteckige weiße und schwarze Kästchen, je nachdem der entsprechende Binomial-

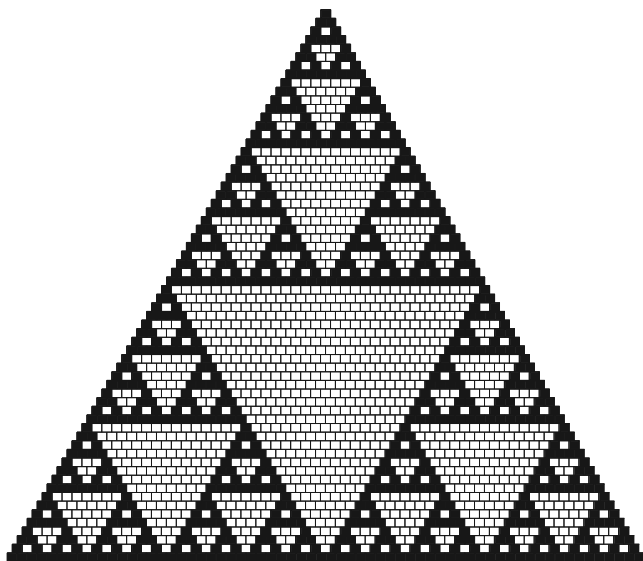


Bild 1.1 Pascalsches Dreieck modulo 2

Koeffizient gerade oder ungerade ist, so entsteht eine interessante Figur, siehe Bild 1.1. Wir bezeichnen das Kästchen, das dem Binomial-Koeffizienten $\binom{k}{\ell}$ entspricht, mit (k, ℓ) . In der Figur sind alle Kästchen (k, ℓ) bis $k = 63$ dargestellt. Man beweise dazu:

a) $\binom{2^n-1}{\ell}$ ist ungerade für alle $0 \leq \ell \leq 2^n - 1$, d.h. die Zeile mit $k = 2^n - 1$ ist vollständig schwarz.

b) $\binom{2^n}{\ell}$ ist gerade für alle $1 \leq \ell \leq 2^n - 1$.

c) $\binom{2^n+\ell}{\ell}$ ist ungerade für alle $0 \leq \ell \leq 2^n - 1$.

d) Das Dreieck mit den Ecken $(0, 0)$, $(2^n - 1, 0)$, $(2^n - 1, 2^n - 1)$ geht durch Verschiebung $(k, \ell) \mapsto (2^n + k, \ell)$ in das Dreieck $(2^n, 0)$, $(2^{2n} - 1, 0)$, $(2^{2n} - 1, 2^n - 1)$ mit demselben Farbmuster über.

e) Das Dreieck mit den Ecken $(0, 0)$, $(2^n - 1, 0)$, $(2^n - 1, 2^n - 1)$ weist außerdem eine Symmetrie bzgl. Drehungen um den Mittelpunkt mit Winkel 120 Grad und 240 Grad auf, genauer: Durch die Transformation

$$(k, \ell) \mapsto (2^n - 1 - \ell, k - \ell), \quad (0 \leq \ell \leq k \leq 2^n - 1)$$

geht das Dreieck unter Erhaltung des Farbmusters in sich über, d.h. die Binomial-Koeffizienten

$$\binom{k}{\ell} \quad \text{und} \quad \binom{2^n - 1 - \ell}{k - \ell}$$

sind entweder beide gerade oder beide ungerade.

1.13. Sei n eine natürliche Zahl. Wieviele Tripel (k_1, k_2, k_3) natürlicher Zahlen gibt es, die

$$k_1 + k_2 + k_3 = n$$

erfüllen?

1.14. Wie groß ist die Wahrscheinlichkeit, dass beim Lotto „6 aus 49“ alle 6 gezogenen Zahlen gerade (bzw. alle ungerade) sind?

1.15. Es werde zufällig eine 7-stellige Zahl gewählt, wobei jede Zahl von 1 000 000 bis 9 999 999 mit der gleichen Wahrscheinlichkeit auftrete. Wie groß ist die Wahrscheinlichkeit dafür, dass alle 7 Ziffern paarweise verschieden sind?

1.16. Man zeige, dass nach dem Gregorianischen Kalender (d.h. Schaltjahr, wenn die Jahreszahl durch 4 teilbar ist, mit Ausnahme der Jahre, die durch 100 aber nicht durch 400 teilbar sind) der 13. eines Monats im langjährigen Durchschnitt häufiger auf einen Freitag fällt, als auf irgend einen anderen Wochentag. Hinweis: Der Geburtstag von Gauß, der 30. April 1777, war ein Mittwoch. (Diese Aufgabe ist weniger eine Übung zur vollständigen Induktion, als eine Übung im systematischen Abzählen.)

§ 2 Die Körper-Axiome

Wir setzen in diesem Buch die reellen Zahlen als gegeben voraus. Um auf sicherem Boden zu stehen, werden wir in diesem und den folgenden Paragraphen einige Axiome formulieren, aus denen sich alle Eigenschaften und Gesetze der reellen Zahlen ableiten lassen.

In diesem Paragraphen behandeln wir die sogenannten Körper-Axiome, aus denen die Rechenregeln für die vier Grundrechnungsarten folgen. Da diese Rechenregeln sämtlich aus dem Schulunterricht geläufig sind, und dem Anfänger erfahrungsgemäß Beweise selbstverständlich erscheinender Aussagen Schwierigkeiten machen, kann dieser Paragraph bei der ersten Lektüre übergangen werden.

Mit $\mathbb{R}$ sei die Menge aller reellen Zahlen bezeichnet. Auf $\mathbb{R}$ sind zwei Verknüpfungen (Addition und Multiplikation)

$$\begin{aligned} + : \mathbb{R} \times \mathbb{R} &\longrightarrow \mathbb{R}, & (x, y) &\mapsto x + y, \\ \cdot : \mathbb{R} \times \mathbb{R} &\longrightarrow \mathbb{R}, & (x, y) &\mapsto xy, \end{aligned}$$

gegeben, die den sog. Körper-Axiomen genügen. Diese bestehen aus den Axiomen der Addition, der Multiplikation und dem Distributivgesetz, die wir der Reihe nach besprechen.

I. Axiome der Addition

(A.1) Assoziativgesetz. *Für alle $x, y, z \in \mathbb{R}$ gilt*

$$(x + y) + z = x + (y + z).$$

(A.2) Kommutativgesetz. *Für alle $x, y \in \mathbb{R}$ gilt*

$$x + y = y + x.$$

(A.3) Existenz der Null. *Es gibt eine Zahl $0 \in \mathbb{R}$, so dass*

$$x + 0 = x \quad \text{für alle } x \in \mathbb{R}.$$

(A.4) Existenz des Negativen. *Zu jedem $x \in \mathbb{R}$ existiert eine Zahl $-x \in \mathbb{R}$, so dass*

$$x + (-x) = 0.$$

Folgerungen aus den Axiomen der Addition

(2.1) Die Zahl 0 ist durch ihre Eigenschaft eindeutig bestimmt.

Beweis. Sei $0' \in \mathbb{R}$ ein weiteres Element mit $x + 0' = x$ für alle $x \in \mathbb{R}$. Dann gilt insbesondere $0 + 0' = 0$. Andererseits ist $0' + 0 = 0'$ nach Axiom (A.3). Da nach dem Kommutativgesetz (A.2) aber $0 + 0' = 0' + 0$, folgt $0 = 0'$, q.e.d.

(2.2) Das Negative einer Zahl $x \in \mathbb{R}$ ist eindeutig bestimmt.

Beweis. Sei x' eine reelle Zahl mit $x + x' = 0$. Addition von $-x$ von links auf beiden Seiten der Gleichung ergibt $(-x) + (x + x') = (-x) + 0$. Nach den Axiomen (A.1) und (A.3) folgt daraus

$$((-x) + x) + x' = -x.$$

Nach (A.2) und (A.4) ist $(-x) + x = x + (-x) = 0$, also

$$((-x) + x) + x' = 0 + x' = x' + 0 = x'.$$

Durch Vergleich erhält man $-x = x'$, q.e.d.

(2.3) Es gilt $-0 = 0$.

Beweis. Nach (A.4) gilt $0 + (-0) = 0$ und nach (A.3) ist $0 + 0 = 0$. Da aber das Negative von 0 eindeutig bestimmt ist, folgt $-0 = 0$.

Bezeichnung. Für $x, y \in \mathbb{R}$ setzt man $x - y := x + (-y)$.

(2.4) Die Gleichung $a + x = b$ hat eine eindeutig bestimmte Lösung, nämlich $x = b - a$.

Beweis. i) Wir zeigen zunächst, dass $x = b - a$ die Gleichung löst. Es ist nämlich

$$\begin{aligned} a + (b - a) &= a + (b + (-a)) = a + ((-a) + b) = (a + (-a)) + b \\ &= 0 + b = b + 0 = b, \quad \text{q.e.d.} \end{aligned}$$

Dabei wurden bei den Umformungen die Axiome (A.1) bis (A.4) benutzt.

ii) Wir zeigen jetzt die Eindeutigkeit der Lösung. Sei y irgend eine Zahl mit $a + y = b$. Addition von $-a$ auf beiden Seiten ergibt

$$(-a) + (a + y) = (-a) + b.$$

Die linke Seite der Gleichung ist gleich $((-a) + a) + y = 0 + y = y$, die rechte Seite gleich $b + (-a) = b - a$, d.h. es gilt $y = b - a$, q.e.d.

(2.5) Für jedes $x \in \mathbb{R}$ gilt $-(-x) = x$.

Beweis. Nach Definition des Negativen von $-x$ gilt $(-x) + (-(-x)) = 0$. Andererseits ist nach (A.2) und (A.4) auch $(-x) + x = x + (-x) = 0$. Aus der Eindeutigkeit des Negativen folgt nun $-(-x) = x$.

(2.6) Für alle $x, y \in \mathbb{R}$ gilt $-(x+y) = -x-y$.

Beweis. Nach Definition des Negativen von $x+y$ ist $(x+y) + (-(x+y)) = 0$. Addition von $-x$ auf beiden Seiten der Gleichung liefert

$$y + (-(x+y)) = -x.$$

Andererseits hat die Gleichung $y+z = -x$ für z die eindeutig bestimmte Lösung $z = -x-y$. Daraus folgt $-(x+y) = -x-y$, q.e.d.

II. Axiome der Multiplikation

(M.1) Assoziativgesetz. Für alle $x, y, z \in \mathbb{R}$ gilt

$$(xy)z = x(yz).$$

(M.2) Kommutativgesetz. Für alle $x, y \in \mathbb{R}$ gilt

$$xy = yx.$$

(M.3) Existenz der Eins. Es gibt ein Element $1 \in \mathbb{R}$, $1 \neq 0$, so dass

$$x \cdot 1 = x \quad \text{für alle } x \in \mathbb{R}.$$

(M.4) Existenz des Inversen. Zu jedem $x \in \mathbb{R}$ mit $x \neq 0$ gibt es ein $x^{-1} \in \mathbb{R}$, so dass

$$xx^{-1} = 1.$$

III. Distributivgesetz

(D) Für alle $x, y, z \in \mathbb{R}$ gilt $x(y+z) = xy+xz$.

Folgerungen aus den Axiomen II und III

(2.7) Die Eins ist durch ihre Eigenschaft eindeutig bestimmt.

(2.8) Das Inverse einer reellen Zahl $x \neq 0$ ist eindeutig bestimmt.

Die Aussagen (2.7) und (2.8) werden ganz analog den entsprechenden Aussagen (2.1) und (2.2) für die Addition bewiesen, indem man überall die Addition

durch die Multiplikation, die Null durch die Eins und das Negative durch das Inverse ersetzt.

(2.9) Für alle $a, b \in \mathbb{R}$ mit $a \neq 0$ hat die Gleichung $ax = b$ eine eindeutig bestimmte Lösung, nämlich $x = a^{-1}b =: \frac{b}{a} =: b/a$.

Beweis. i) $x = a^{-1}b$ löst die Gleichung, denn

$$a(a^{-1}b) = (aa^{-1})b = 1 \cdot b = b \cdot 1 = b.$$

ii) Zur Eindeutigkeit. Sei y eine beliebige Zahl mit $ay = b$. Multiplikation der Gleichung mit a^{-1} von links ergibt $a^{-1}(ay) = a^{-1}b$. Die linke Seite der Gleichung kann man unter Anwendung der Axiome (M.1) bis (M.4) umformen und erhält $a^{-1}(ay) = y$, woraus folgt $y = a^{-1}b$, q.e.d.

(2.10) Für alle $x, y, z \in \mathbb{R}$ gilt $(x + y)z = xz + yz$.

Beweis. Unter Benutzung von (M.2) und (D) erhalten wir

$$(x + y)z = z(x + y) = zx + zy = xz + yz, \quad \text{q.e.d.}$$

(2.11) Für alle $x \in \mathbb{R}$ gilt $x \cdot 0 = 0$.

Beweis. Da $0 + 0 = 0$, folgt aus dem Distributivgesetz

$$x \cdot 0 + x \cdot 0 = x \cdot (0 + 0) = x \cdot 0.$$

Subtraktion von $x \cdot 0$ von beiden Seiten der Gleichung ergibt $x \cdot 0 = 0$.

(2.12) Für $x, y \in \mathbb{R}$ gilt $xy = 0$ genau dann, wenn $x = 0$ oder $y = 0$.

(In Worten: Ein Produkt ist genau dann gleich null, wenn einer der Faktoren null ist.)

Beweis. Wenn $x = 0$ oder $y = 0$, so folgt aus (2.11), dass $xy = 0$. Sei nun umgekehrt vorausgesetzt, dass $xy = 0$. Falls $x = 0$, sind wir fertig. Falls aber $x \neq 0$, folgt aus (2.9), dass $y = x^{-1} \cdot 0 = 0$, q.e.d.

(2.13) Für alle $x \in \mathbb{R}$ gilt $-x = (-1)x$.

Beweis. Unter Benutzung des Distributivgesetzes erhält man

$$x + (-1) \cdot x = 1 \cdot x + (-1) \cdot x = (1 - 1) \cdot x = 0 \cdot x = 0,$$

d.h. $(-1)x$ ist ein Negatives von x . Wegen der Eindeutigkeit des Negativen folgt die Behauptung.

(2.14) Für alle $x, y \in \mathbb{R}$ gilt $(-x)(-y) = xy$.

Beweis. Mit (2.13), sowie dem Kommutativ- und Assoziativgesetz erhält man

$$(-x)(-y) = (-x)(-1)y = (-1)(-x)y = (-(-x))y.$$

Da $-(-x) = x$ wegen (2.5), folgt die Behauptung.

(2.15) Für alle reellen Zahlen $x \neq 0$ gilt $(x^{-1})^{-1} = x$.

(2.16) Für alle reellen Zahlen $x \neq 0, y \neq 0$ gilt $(xy)^{-1} = x^{-1}y^{-1}$.

Die Regeln (2.15) und (2.16) sind die multiplikativen Analoga der Regeln (2.5) und (2.6) und können auch analog bewiesen werden.

Allgemeines Assoziativgesetz

Die Addition von mehr als zwei Zahlen wird durch Klammerung auf die Addition von jeweils zwei Summanden zurückgeführt:

$$x_1 + x_2 + x_3 + \dots + x_n := (\dots((x_1 + x_2) + x_3) + \dots) + x_n.$$

Man beweist durch wiederholte Anwendung des Assoziativgesetzes (A.1), dass jede andere Klammerung zum selben Resultat führt. Analoges gilt für das Produkt $x_1 x_2 \cdot \dots \cdot x_n$.

Allgemeines Kommutativgesetz

Sei $(i_1, i_2, \dots, i_n)$ eine Permutation (d.h. Umordnung) von $(1, 2, \dots, n)$. Dann gilt

$$x_1 + x_2 + \dots + x_n = x_{i_1} + x_{i_2} + \dots + x_{i_n},$$

$$x_1 x_2 \cdot \dots \cdot x_n = x_{i_1} x_{i_2} \cdot \dots \cdot x_{i_n}.$$

Dies folgt durch wiederholte Anwendung der Kommutativgesetze (A.2) bzw. (M.2) sowie der Assoziativgesetze.

Aus dem allgemeinen Kommutativgesetz kann man folgende Regel für *Doppeisummen* ableiten:

$$\sum_{i=1}^n \sum_{j=1}^m a_{ij} = \sum_{j=1}^m \sum_{i=1}^n a_{ij}.$$

Denn nach Definition gilt

$$\sum_{i=1}^n \sum_{j=1}^m a_{ij} = \left(\sum_{j=1}^m a_{1j} \right) + \left(\sum_{j=1}^m a_{2j} \right) + \dots + \left(\sum_{j=1}^m a_{nj} \right)$$

$$= (a_{11} + a_{12} + \dots + a_{1m}) + \dots + (a_{n1} + a_{n2} + \dots + a_{nm})$$

und

$$\begin{aligned} \sum_{j=1}^m \sum_{i=1}^n a_{ij} &= \left(\sum_{i=1}^n a_{i1} \right) + \left(\sum_{i=1}^n a_{i2} \right) + \dots + \left(\sum_{i=1}^n a_{im} \right) \\ &= (a_{11} + a_{21} + \dots + a_{n1}) + \dots + (a_{1m} + a_{2m} + \dots + a_{nm}). \end{aligned}$$

Es kommen also in beiden Fällen alle nm Summanden a_{ij} , $1 \leq i \leq n$, $1 \leq j \leq m$, vor, nur in anderer Reihenfolge.

Allgemeines Distributivgesetz

Durch wiederholte Anwendung von (D) und Folgerung (2.10) beweist man

$$\left(\sum_{i=1}^n x_i \right) \left(\sum_{j=1}^m y_j \right) = \sum_{i=1}^n \sum_{j=1}^m x_i y_j.$$

Potenzen

Ist x eine reelle Zahl, so werden die Potenzen x^n für $n \in \mathbb{N}$ durch Induktion wie folgt definiert:

$$x^0 := 1, \quad x^{n+1} := x^n x \quad \text{für alle } n \geq 0.$$

(Man beachte, dass nach Definition auch $0^0 = 1$.)

Ist $x \neq 0$, so definiert man negative Potenzen x^{-n} , ($n > 0$ ganz), durch

$$x^{-n} := (x^{-1})^n.$$

Für die Potenzen gelten folgende Rechenregeln:

$$(2.17) \quad x^n x^m = x^{n+m},$$

$$(2.18) \quad (x^n)^m = x^{nm},$$

$$(2.19) \quad x^n y^n = (xy)^n.$$

Dabei sind n und m beliebige ganze Zahlen und x, y reelle Zahlen, die $\neq 0$ vorauszusetzen sind, falls negative Exponenten vorkommen.

Wir beweisen als Beispiel die Aussage (2.19) und überlassen die anderen dem Leser als Übung.

a) Falls $n \geq 0$, verwenden wir vollständige Induktion nach n . Der Induktions-Anfang $n = 0$ ist trivial.

Induktions-Schritt $n \rightarrow n+1$. Unter Verwendung des Kommutativ- und Assoziativgesetzes der Multiplikation erhält man

$$x^{n+1}y^{n+1} = x^n xy^n y = x^n y^n xy \stackrel{\text{IV}}{=} (xy)^n xy = (xy)^{n+1}, \quad \text{q.e.d.}$$

b) Falls $n < 0$, ist $m := -n > 0$ und

$$x^n y^n = x^{-m} y^{-m} = (x^{-1})^m (y^{-1})^m.$$

Nach a) gilt $(x^{-1})^m (y^{-1})^m = (x^{-1} y^{-1})^m$, also unter Benutzung von (2.16)

$$x^n y^n = (x^{-1} y^{-1})^m = ((xy)^{-1})^m = (xy)^{-m} = (xy)^n, \quad \text{q.e.d.}$$

Bemerkung. Eine Menge K , zusammen mit zwei Verknüpfungen

$$+ : K \times K \longrightarrow K, \quad (x, y) \mapsto x + y,$$

$$\cdot : K \times K \longrightarrow K, \quad (x, y) \mapsto xy,$$

die den Axiomen I bis III genügen, nennt man *Körper*. In jedem Körper gelten alle in diesem Paragraphen hergeleiteten Rechenregeln, da zu ihrem Beweis nur die Axiome verwendet wurden.

Beispiele. $\mathbb{R}$, $\mathbb{Q}$ (Menge der rationalen Zahlen), und $\mathbb{C}$ (Menge der komplexen Zahlen, siehe § 13) bilden mit der üblichen Addition und Multiplikation jeweils einen Körper. Dagegen ist die Menge $\mathbb{Z}$ aller ganzen Zahlen kein Körper, da das Axiom von der Existenz des Inversen verletzt ist (z.B. besitzt die Zahl $2 \in \mathbb{Z}$ in $\mathbb{Z}$ kein Inverses).

Ein merkwürdiger Körper ist die Menge $\mathbb{F}_2 = \{0, 1\}$ mit den Verknüpfungen

$$\begin{array}{c|cc} + & 0 & 1 \\ \hline 0 & 0 & 1 \\ 1 & 1 & 0 \end{array} \quad \text{und} \quad \begin{array}{c|cc} \cdot & 0 & 1 \\ \hline 0 & 0 & 0 \\ 1 & 0 & 1 \end{array}$$

Die Körper-Axiome können hier durch direktes Nachprüfen aller Fälle verifiziert werden. $\mathbb{F}_2$ ist der kleinst-mögliche Körper, denn jeder Körper muss mindestens die Null und die Eins enthalten. In $\mathbb{F}_2$ gilt $1 + 1 = 0$. Also kann man die Aussage $1 + 1 \neq 0$ nicht mithilfe der Körper-Axiome beweisen. Insbesondere kann man allein aufgrund der Körper-Axiome die natürlichen Zahlen noch nicht als Teilmenge der reellen Zahlen auffassen. Hierzu sind weitere Axiome erforderlich, die wir im nächsten Paragraphen behandeln.

AUFGABEN

2.1. Man zeige: Es gelten die folgenden Regeln für das Bruchrechnen ($a, b, c, d \in \mathbb{R}, b \neq 0, d \neq 0$):

a) $\frac{a}{b} = \frac{c}{d}$ genau dann, wenn $ad = bc$,

b) $\frac{a}{b} \pm \frac{c}{d} = \frac{ad \pm bc}{bd}$,

c) $\frac{a}{b} \cdot \frac{c}{d} = \frac{ac}{bd}$,

d) $\frac{\frac{a}{b}}{\frac{c}{d}} = \frac{ad}{bc}$, falls $c \neq 0$.

2.2. Man beweise die Rechenregel (2.17) für Potenzen:

$$x^n x^m = x^{n+m}, \quad (n, m \in \mathbb{Z}, x \in \mathbb{R}, \text{ wobei } x \neq 0 \text{ falls } n < 0 \text{ oder } m < 0).$$

Anleitung. Man behandle zunächst die Fälle

(1) $n \geq 0, m \geq 0$,

(2) $n > 0$ und $m = -k$ mit $0 < k \leq n$,

und führe den allgemeinen Fall auf (1) und (2) zurück.

2.3. Seien a_{ik} für $i, k \in \mathbb{N}$ reelle Zahlen. Man zeige für alle $n \in \mathbb{N}$

$$\sum_{k=0}^n \sum_{i=0}^{n-k} a_{ik} = \sum_{i=0}^n \sum_{k=0}^{n-i} a_{ik} = \sum_{m=0}^n \sum_{k=0}^m a_{m-k,k}.$$

2.4. Es sei $\overline{\mathbb{N}} := \mathbb{N} \cup \{\infty\}$, wobei ∞ ein nicht zu $\mathbb{N}$ gehöriges Symbol ist. Auf $\overline{\mathbb{N}}$ führen wir zwei Verknüpfungen

$$\overline{\mathbb{N}} \times \overline{\mathbb{N}} \rightarrow \overline{\mathbb{N}}, \quad (a, b) \mapsto a + b,$$

$$\overline{\mathbb{N}} \times \overline{\mathbb{N}} \rightarrow \overline{\mathbb{N}}, \quad (a, b) \mapsto a \cdot b,$$

wie folgt ein:

i) Für $a, b \in \mathbb{N}$ sei $a + b$ bzw. $a \cdot b$ die übliche Addition bzw. Multiplikation natürlicher Zahlen.

ii) $a + \infty = \infty + a = \infty$ für alle $a \in \overline{\mathbb{N}}$.

iii) $0 \cdot \infty = \infty \cdot 0 = 0$ und $a \cdot \infty = \infty \cdot a = \infty$ für alle $a \in \overline{\mathbb{N}} \setminus \{0\}$.

Man zeige, dass diese Verknüpfungen auf $\overline{\mathbb{N}}$ die Körperaxiome (A.1), (A.2), (A.3), (M.1), (M.2), (M.3) und (D), aber nicht (A.4) und (M.4) erfüllen.

§ 3 Die Anordnungs-Axiome

In der Analysis ist das Rechnen mit Ungleichungen ebenso wichtig wie das Rechnen mit Gleichungen. Das Rechnen mit Ungleichungen beruht auf den Anordnungs-Axiomen. Es stellt sich heraus, dass alles auf den Begriff des positiven Elements zurückgeführt werden kann.

Anordnungs-Axiome. In $\mathbb{R}$ sind gewisse Elemente als positiv ausgezeichnet (Schreibweise $x > 0$), so dass folgende Axiome erfüllt sind.

(O.1) Trichotomie. Für jedes x gilt genau eine der drei Beziehungen

$$x > 0, \quad x = 0, \quad -x > 0.$$

(O.2) Abgeschlossenheit gegenüber Addition.

$$x > 0 \text{ und } y > 0 \implies x + y > 0.$$

(O.3) Abgeschlossenheit gegenüber Multiplikation.

$$x > 0 \text{ und } y > 0 \implies xy > 0.$$

Die Axiome (O.2) und (O.3) lassen sich zusammenfassend kurz so ausdrücken: Summe und Produkt positiver Elemente sind wieder positiv.

Zur Notation. Wir haben hier in der Formulierung der Axiome den Implikationspfeil benutzt. $A \Rightarrow B$ bedeutet, dass die Aussage B aus der Aussage A folgt. Die Bezeichnung $A \Leftrightarrow B$ bedeutet, dass sowohl $A \Rightarrow B$ als auch $B \Rightarrow A$ gilt, also die Aussagen A und B logisch äquivalent sind. Schließlich heißt die Bezeichnung $A :\Leftrightarrow B$, dass die Aussage A durch die Aussage B definiert wird.

Definition (Größer- und Kleiner-Relation). Für reelle Zahlen x, y definiert man

$$\begin{aligned} x > y & :\Leftrightarrow x - y > 0, \\ x < y & :\Leftrightarrow y > x, \\ x \geq y & :\Leftrightarrow x > y \text{ oder } x = y, \\ x \leq y & :\Leftrightarrow x < y \text{ oder } x = y. \end{aligned}$$

Folgerungen aus den Axiomen

In den folgenden Aussagen sind x, y, z, a, b stets Elemente von $\mathbb{R}$.

(3.1) Für zwei Elemente x, y gilt genau eine der Relationen

$$x < y, \quad x = y, \quad y < x.$$

Dies folgt unmittelbar aus Axiom (O.1). Damit kann man das Maximum und Minimum zweier reeller Zahlen definieren:

$$\max(x, y) := \begin{cases} x, & \text{falls } x \geq y, \\ y & \text{sonst,} \end{cases}$$

$$\min(x, y) := \begin{cases} x, & \text{falls } x \leq y, \\ y & \text{sonst.} \end{cases}$$

(3.2) *Transitivität der Kleiner-Relation*

$$x < y \text{ und } y < z \implies x < z$$

Beweis. Die Voraussetzungen bedeuten $y - x > 0$ und $z - y > 0$. Mit Axiom (O.2) folgt daraus $(y - x) + (z - y) = z - x > 0$, d.h. $x < z$.

(3.3) *Translations-Invarianz*

$$x < y \implies a + x < a + y$$

Dies folgt nach Definition daraus, dass $(a + y) - (a + x) = y - x$.

(3.4) *Spiegelung*

$$x < y \iff -x > -y$$

Dies folgt aus $y - x = (-x) - (-y)$.

Diese Aussagen unterstützen unsere anschauliche Vorstellung der reellen Zahlengeraden. Zeichnet man die Zahlengerade waagrecht, so denkt man sich die positiven Zahlen rechts vom Nullpunkt, die negativen Zahlen links davon. Von zwei Zahlen ist diejenige die größere, die weiter rechts liegt. Addition einer Zahl a entspricht einer Verschiebung (nach rechts, wenn $a > 0$, nach links, wenn $a < 0$). Der Übergang von x zu $-x$ bedeutet eine Spiegelung am Nullpunkt; dabei werden die Rollen von rechts und links vertauscht.

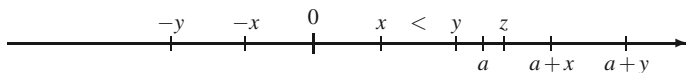


Bild 3.1 Die Zahlengerade

(Es ist natürlich nur eine Konvention, dass die Zahlen in Richtung von links nach rechts größer werden; man hätte genauso gut die andere Richtung wählen können. Die übliche Konvention erklärt sich wohl aus der Schreibrichtung von links nach rechts.)

$$(3.5) \quad x < y \text{ und } a < b \implies x + a < y + b$$

Beweis. Mit (3.3) folgt aus den Voraussetzungen $a + x < a + y$ und $y + a < y + b$. Wegen der Transitivität ergibt sich daraus $x + a < y + b$.

$$(3.6) \quad x < y \text{ und } a > 0 \implies ax < ay$$

Kurz gesagt: Man darf eine Ungleichung mit einer positiven Zahl multiplizieren.

Beweis. Da nach Voraussetzung $y - x > 0$ und $a > 0$, folgt aus Axiom (O.3), dass $a(y - x) = ay - ax > 0$. Dies bedeutet aber nach Definition $ax < ay$.

$$(3.7) \quad 0 \leq x < y \text{ und } 0 \leq a < b \implies ax < by$$

Beweis. Steht bei einer der beiden Voraussetzungen das Gleichheitszeichen, so ist stets $ax = 0 < by$. Sei also $0 < x < y$ und $0 < a < b$. Mit (3.6) folgt $ax < ay$ und $ay < by$, also aufgrund der Transitivität $ax < by$.

$$(3.8) \quad x < y \text{ und } a < 0 \implies ax > ay$$

Anders ausgedrückt: Multipliziert man eine Ungleichung mit einer negativen Zahl, so verwandelt sich das Kleiner- in ein Größer-Zeichen.

Beweis. Da $-a > 0$ (nach (3.4)), erhält man mit (3.6) $-ax < -ay$. Die Behauptung folgt durch nochmalige Anwendung von (3.4).

(3.9) Für jedes Element $x \neq 0$ ist $x^2 > 0$, insbesondere gilt $1 > 0$.

Beweis. Ist $x > 0$, so folgt $x^2 > 0$ aus Axiom (O.3); ist dagegen $x < 0$, so folgt dies aus (3.8). Da $0 \neq 1 = 1^2$, ergibt sich $1 > 0$.

$$(3.10) \quad x > 0 \iff x^{-1} > 0$$

Beweis. Da $x^{-2} > 0$ nach (3.9), ergibt sich die Implikation ' $\Rightarrow$ ' durch Multiplikation von x mit x^{-2} aus Axiom (O.3). Die Umkehrung ' $\Leftarrow$ ' folgt aus ' $\Rightarrow$ ', angewendet auf x^{-1} , da $(x^{-1})^{-1} = x$.

$$(3.11) \quad 0 < x < y \implies x^{-1} > y^{-1}$$

Beweis. Aus Axiom (O.3) folgt $xy > 0$, also nach (3.10) auch $(xy)^{-1} = x^{-1}y^{-1} > 0$. Nach (3.6) darf man die Ungleichung $x < y$ mit $x^{-1}y^{-1}$ multiplizieren und erhält

$$y^{-1} = x(x^{-1}y^{-1}) < y(x^{-1}y^{-1}) = x^{-1}, \quad \text{q.e.d.}$$

Bemerkung. Ein Körper, in dem gewisse Elemente als positiv ausgezeichnet sind, so dass die Axiome (O.1), (O.2) und (O.3) gelten, heißt *angeordneter Körper*. $\mathbb{R}$ und $\mathbb{Q}$ sind angeordnete Körper. Dagegen kann der Körper $\mathbb{F}_2$ nicht angeordnet werden, denn in ihm gilt $1 + 1 = 0$, was wegen (3.9) im Widerspruch zu Axiom (O.2) steht. Ebenso besitzt der Körper der komplexen Zahlen (den wir in §13 einführen), keine Anordnung, da in ihm $i^2 = -1$, was der Regel (3.9) widerspricht.

Die natürlichen Zahlen als Teilmenge von $\mathbb{R}$

In jedem Körper gibt es die 0 und die 1. Um die weiteren natürlichen Zahlen zu erhalten, kann man versuchen, einfach sukzessive die 1 zu addieren, $2 := 1+1$, $3 := 2+1$, $4 := 3+1$, usw. Dass dies nicht ohne weiteres das Erwartete liefert, sieht man am Körper $\mathbb{F}_2$, in dem damit $2 = 0$ ist, was unseren Vorstellungen von den natürlichen Zahlen widerspricht. Es stellt sich aber heraus, dass aufgrund der Anordnungs-Axiome innerhalb des Körpers der reellen Zahlen solche Pathologien nicht auftreten können.

Es sei $\mathcal{N}$ die kleinste Teilmenge von $\mathbb{R}$ mit folgenden Eigenschaften:

- i) $0 \in \mathcal{N}$,
- ii) $x \in \mathcal{N} \Rightarrow x + 1 \in \mathcal{N}$.

$\mathcal{N}$ besteht also genau aus den Zahlen, die sich aus der 0 durch sukzessive Additionen von 1 erhalten lassen. Wir definieren eine Abbildung (Nachfolger-Funktion)

$$v: \mathcal{N} \longrightarrow \mathcal{N}, \quad v(x) := x + 1.$$

Um zu sehen, dass die Menge $\mathcal{N}$ aus den ‘richtigen’ natürlichen Zahlen besteht, verifizieren wir die sog. *Peano-Axiome*. Nach Peano lassen sich die natürlichen Zahlen charakterisieren als eine Menge $\mathcal{N}$ mit einem ausgezeichneten Element 0 und einer Abbildung $v: \mathcal{N} \rightarrow \mathcal{N}$, so dass folgende Axiome erfüllt sind:

(P.1) $x \neq y \implies v(x) \neq v(y)$, d.h. zwei verschiedene Elemente von $\mathcal{N}$ haben auch verschiedene Nachfolger.

(P.2) $0 \notin v(\mathcal{N})$, d.h. kein Element von $\mathcal{N}$ hat 0 als Nachfolger.

(P.3) (Induktions-Axiom)

Sei $M \subset \mathcal{N}$ eine Teilmenge mit folgenden Eigenschaften:

i) $0 \in M$,

ii) $x \in M \implies v(x) \in M$.

Dann gilt $M = \mathcal{N}$.

Für unsere Menge $\mathcal{N} \subset \mathbb{R}$ ist (P.3) nach Definition erfüllt und (P.1) ist trivial, denn in jedem Körper folgt aus $x + 1 = y + 1$, dass $x = y$. Es bleibt also nur noch (P.2) nachzuprüfen. Dazu zeigen wir zunächst

$$x \geq 0 \quad \text{für alle } x \in \mathcal{N}.$$

Beweis hierfür. Wir definieren $M := \{x \in \mathcal{N} : x \geq 0\}$. Offenbar erfüllt M die Bedingungen i) und ii) von (P.3), also muss $M = \mathcal{N}$ sein, q.e.d.

Wäre nun (P.2) falsch, so gäbe es ein $x \in \mathcal{N}$ mit $0 = v(x) = x + 1$, also $x = -1$. Nach (3.9) und (3.4) ist $-1 < 0$, im Widerspruch zu $x \geq 0$.

Somit sind alle Peano-Axiome erfüllt. Man kann zeigen, dass durch die Peano-Axiome die natürlichen Zahlen bis auf Isomorphie eindeutig festgelegt sind. Wir werden deshalb $\mathcal{N}$ und $\mathbb{N}$ identifizieren.

Übrigens enthält das Peano-Axiom (P.3) das Prinzip der vollständigen Induktion. Denn sei $A(n)$ für jedes $n \in \mathbb{N}$ eine Aussage. Wir definieren M als die Menge aller $n \in \mathbb{N}$, für die $A(n)$ wahr ist. Dann bedeutet (P.3.i) den Induktions-Anfang und (P.3.ii) den Induktions-Schritt. Sind beide erfüllt, so gilt $M = \mathbb{N}$, d.h. $A(n)$ ist wahr für alle $n \in \mathbb{N}$.

Bemerkung. Die gemachten Überlegungen zeigen, dass die natürlichen Zahlen in jedem angeordneten Körper enthalten sind.

Der Absolut-Betrag

Für eine reelle Zahl x wird ihr (Absolut-)Betrag definiert durch

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

gesprochen *x-Betrag* oder *x-absolut*. Die Definition ist gleichwertig mit

$$|x| := \max(x, -x).$$

Satz 1. *Der Absolut-Betrag in $\mathbb{R}$ hat folgende Eigenschaften:*

a) *Es ist $|x| \geq 0$ für alle $x \in \mathbb{R}$ und*

$$|x| = 0 \iff x = 0.$$

b) (Multiplikativität)

$$|xy| = |x| \cdot |y| \quad \text{für alle } x, y \in \mathbb{R}.$$

c) (Dreiecks-Ungleichung)

$$|x + y| \leq |x| + |y| \quad \text{für alle } x, y \in \mathbb{R}.$$

Beweis. Die Eigenschaft a) folgt unmittelbar aus der Definition.

b) Die Aussage ist trivial für $x, y \geq 0$. Im allgemeinen Fall schreiben wir $x = \pm x_0$ und $y = \pm y_0$ mit $x_0, y_0 \geq 0$. Dann ist

$$|xy| = |\pm x_0 y_0| = |x_0 y_0| = |x_0| \cdot |y_0| = |x| \cdot |y|, \quad \text{q.e.d.}$$

c) Da $x \leq |x|$ und $y \leq |y|$, folgt aus (3.3) und (3.5), dass

$$x + y \leq |x| + |y|.$$

Ebenso ist wegen $-x \leq |x|$ und $-y \leq |y|$

$$-(x + y) = -x - y \leq |x| + |y|.$$

Zusammen genommen ergibt sich $|x + y| \leq |x| + |y|$.

Bemerkung. Ein Körper K , auf dem eine Abbildung

$$K \rightarrow \mathbb{R}, \quad x \mapsto |x|,$$

definiert ist, so dass die in Satz 1 genannten Eigenschaften a), b), c) erfüllt sind, heißt *bewerteter Körper*. Es gibt auch nicht angeordnete bewertete Körper, wie den Körper der komplexen Zahlen, den wir in §13 untersuchen werden (dort erklärt sich auch der Name Dreiecks-Ungleichung). Bei der folgenden Ableitung weiterer Eigenschaften des Absolut-Betrages verwenden wir nur die Regeln a) bis c); sie sind damit in jedem bewerteten Körper gültig.

(3.12) Setzt man in b) $x = y = 1$, erhält man $|1| = |1||1|$, woraus folgt $|1| = 1$. Für $x = y = -1$ ergibt sich $|-1|^2 = |1| = 1$, also $|-1| = 1$ wegen a). Daraus folgt

$$|-x| = |x| \quad \text{für alle } x.$$

(3.13) Für alle $x, y \in \mathbb{R}$ mit $y \neq 0$ gilt

$$\left| \frac{x}{y} \right| = \frac{|x|}{|y|}.$$

Beweis. Weil $x = \frac{x}{y} \cdot y$, folgt aus der Multiplikativität des Betrages $|x| = \left| \frac{x}{y} \right| \cdot |y|$. Bringt man $|y|$ auf die andere Seite (d.h. multipliziert man die Gleichung mit $\frac{1}{|y|}$), erhält man die Behauptung.

(3.14) Für alle $x, y \in \mathbb{R}$ gilt

$$|x - y| \geq |x| - |y| \quad \text{und} \quad |x + y| \geq |x| - |y|.$$

Beweis. Aus $x = (x - y) + y$ erhält man mit der Dreiecks-Ungleichung $|x| \leq |x - y| + |y|$. Addition von $-|y|$ auf beiden Seiten ergibt die erste Ungleichung. Ersetzt man y durch $-y$, folgt daraus die zweite.

Das Archimedische Axiom

Wir benötigen noch ein weiteres sich auf die Anordnung beziehendes Axiom:

(Arch) Zu je zwei reellen Zahlen $x, y > 0$ existiert eine natürliche Zahl n mit $nx > y$.

Archimedes hat dieses Axiom im Rahmen der Geometrie formuliert: Hat man zwei Strecken auf einer Geraden, so kann man, wenn man die kleinere von beiden nur oft genug abträgt, die größere übertreffen, siehe dazu Bild 3.2.



Bild 3.2 Zum Archimedischen Axiom

Bemerkung. Ein angeordneter Körper, in dem das Archimedische Axiom gilt, heißt *archimedisch angeordnet*. $\mathbb{R}$ und $\mathbb{Q}$ sind archimedisch angeordnete Körper. Es gibt aber angeordnete Körper, in denen das Archimedische Axiom nicht

gilt (siehe z.B. [H]). Also ist das Archimedische Axiom von den bisherigen Axiomen unabhängig.

Folgerungen aus dem Archimedischen Axiom

(3.15) Zu jeder reellen Zahl x gibt es natürliche Zahlen n_1 und n_2 , so dass $n_1 > x$ und $-n_2 < x$. Daraus folgt: Zu jedem $x \in \mathbb{R}$ gibt es eine eindeutig bestimmte ganze Zahl $n \in \mathbb{Z}$ mit

$$n \leq x < n + 1.$$

Diese ganze Zahl wird mit $\lfloor x \rfloor$ oder $\text{floor}(x)$ bezeichnet. Statt $\lfloor x \rfloor$ ist auch die Bezeichnung $[x]$ üblich (Gauß-Klammer).

Ebenso existiert eine eindeutig bestimmte ganze Zahl $m \in \mathbb{Z}$ mit

$$m - 1 < x \leq m,$$

welche mit $\lceil x \rceil$ oder $\text{ceil}(x)$ bezeichnet wird (von engl. *ceiling* = Decke).

(3.16) Zu jedem $\varepsilon > 0$ existiert eine natürliche Zahl $n > 0$ mit

$$\frac{1}{n} < \varepsilon.$$

Beweis. Nach (3.15) existiert ein n mit $n > 1/\varepsilon$. Mit (3.11) folgt daraus $1/n < \varepsilon$.

Satz 2 (Bernoullische Ungleichung). *Sei $x \geq -1$. Dann gilt*

$$(1+x)^n \geq 1+nx \quad \text{für alle } n \in \mathbb{N}.$$

Beweis durch vollständige Induktion nach n .

Induktions-Anfang $n = 0$. Trivialerweise ist $(1+x)^0 = 1 \geq 1$.

Induktions-Schritt $n \rightarrow n+1$. Da $1+x \geq 0$, folgt durch Multiplikation der Induktions-Voraussetzung $(1+x)^n \geq 1+nx$ mit $1+x$

$$\begin{aligned} (1+x)^{n+1} &\geq (1+nx)(1+x) \\ &= 1 + (n+1)x + nx^2 \geq 1 + (n+1)x, \quad \text{q.e.d.} \end{aligned}$$

Satz 3 (Wachstum von Potenzen). *Sei b eine positive reelle Zahl.*

a) *Ist $b > 1$, so gibt es zu jedem $K \in \mathbb{R}$ ein $n \in \mathbb{N}$, so dass*

$$b^n > K.$$

b) *Ist $0 < b < 1$, so gibt es zu jedem $\varepsilon > 0$ ein $n \in \mathbb{N}$, so dass*

$$b^n < \varepsilon.$$

Beweis. a) Sei $x := b - 1$. Nach Voraussetzung ist $x > 0$. Die Bernoullische Ungleichung sagt

$$b^n = (1+x)^n \geq 1+nx.$$

Nach dem Archimedischem Axiom gibt es ein $n \in \mathbb{N}$ mit $nx > K - 1$. Für dieses n ist dann $b^n > K$.

b) Da $b_1 := 1/b > 1$, gibt es nach Teil a) zu $K := 1/\varepsilon$ ein n mit $b_1^n > 1/\varepsilon$. Mit (3.11) folgt $b^n < \varepsilon$, q.e.d.

AUFGABEN

3.1. Man zeige $n^2 \leq 2^n$ für jede natürliche Zahl $n \neq 3$.

3.2. Man zeige $2^n < n!$ für jede natürliche Zahl $n \geq 4$.

3.3. Man beweise:

a) Für jede reelle Zahl x mit $0 \leq x < 1$ gilt

$$\frac{1}{1-x} \geq 1+x.$$

b) Für jede reelle Zahl x mit $0 \leq x \leq \frac{1}{2}$ gilt

$$\frac{1}{1-x} \leq 1+2x.$$

3.4. Man beweise mit Hilfe des Binomischen Lehrsatzes: Für jede reelle Zahl $x \geq 0$ und jede natürliche Zahl $n \geq 2$ gilt

$$(1+x)^n > \frac{n^2}{4}x^2.$$

3.5. Seien $a_1, \dots, a_n$ nicht-negative reelle Zahlen. Man beweise:

$$\prod_{i=1}^n (1 + a_i) \geq 1 + \sum_{i=1}^n a_i.$$

3.6. Seien $a_1, \dots, a_n$ nicht-negative reelle Zahlen mit $\sum_{i=1}^n a_i \leq 1$. Man beweise:

a)
$$\prod_{i=1}^n (1 + a_i) \leq 1 + 2 \sum_{i=1}^n a_i,$$

b)
$$\prod_{i=1}^n (1 - a_i) \geq 1 - \sum_{i=1}^n a_i.$$

3.7. Man beweise: Für jede natürliche Zahl $n \geq 1$ gelten die folgenden Aussagen:

a)
$$\binom{n}{k} \frac{1}{n^k} \leq \frac{1}{k!} \quad \text{für alle } k \in \mathbb{N},$$

b)
$$\left(1 + \frac{1}{n}\right)^n \leq \sum_{k=0}^n \frac{1}{k!} < 3,$$

c)
$$\left(\frac{n}{3}\right)^n \leq \frac{1}{3} n!$$

3.8. Man beweise: Für alle natürlichen Zahlen n gilt

$$n! \leq 2 \left(\frac{n}{2}\right)^n.$$

3.9. Man zeige: Für alle reellen Zahlen x, y gilt

$$\max(x, y) = \frac{1}{2}(x + y + |x - y|),$$

$$\min(x, y) = \frac{1}{2}(x + y - |x - y|).$$

3.10. Man beweise folgende Regeln für die Funktionen floor und ceil:

a)
$$\lceil x \rceil = -\lfloor -x \rfloor \quad \text{für alle } x \in \mathbb{R}.$$

b)
$$\lceil x \rceil = \lfloor x \rfloor + 1 \quad \text{für alle } x \in \mathbb{R} \setminus \mathbb{Z}.$$

c)
$$\lceil n/k \rceil = \lfloor (n+k-1)/k \rfloor \quad \text{für alle } n, k \in \mathbb{Z} \text{ mit } k \geq 1.$$

§ 4 Folgen, Grenzwerte

Wir kommen jetzt zu einem der zentralen Begriffe der Analysis, dem des Grenzwerts einer Folge. Seine Bedeutung beruht darauf, dass viele Größen nicht durch einen in endlich vielen Schritten exakt berechenbaren Ausdruck gegeben, sondern nur mit beliebiger Genauigkeit approximiert werden können. Eine Zahl mit beliebiger Genauigkeit approximieren heißt, sie als Grenzwert einer Folge darstellen. Dies werden wir jetzt präzisieren.

Unter einer *Folge* reeller Zahlen versteht man eine Abbildung $\mathbb{N} \rightarrow \mathbb{R}$. Jedem $n \in \mathbb{N}$ ist also ein $a_n \in \mathbb{R}$ zugeordnet. Man schreibt hierfür

$$(a_n)_{n \in \mathbb{N}} \quad \text{oder} \quad (a_0, a_1, a_2, a_3, \dots)$$

oder kurz (a_n) . Etwas allgemeiner kann man als Indexmenge statt $\mathbb{N}$ die Menge $\{n \in \mathbb{Z} : n \geq k\}$ aller ganzen Zahlen, die größer-gleich einer vorgegebenen ganzen Zahl k sind, zulassen. So erhält man Folgen

$$(a_n)_{n \geq k} \quad \text{oder} \quad (a_k, a_{k+1}, a_{k+2}, \dots).$$

Beispiele

(4.1) Sei $a_n = a$ für alle $n \in \mathbb{N}$. Man erhält die *konstante Folge*

$$(a, a, a, a, \dots).$$

(4.2) Sei $a_n = \frac{1}{n}$, $n \geq 1$. Dies ergibt die Folge

$$(1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots).$$

(4.3) Für $a_n = (-1)^n$ ist

$$(a_n)_{n \in \mathbb{N}} = (+1, -1, +1, -1, +1, \dots).$$

(4.4) $\left(\frac{n}{n+1}\right)_{n \in \mathbb{N}} = (0, \frac{1}{2}, \frac{2}{3}, \frac{3}{4}, \frac{4}{5}, \dots).$

(4.5) $\left(\frac{n}{2^n}\right)_{n \in \mathbb{N}} = (0, \frac{1}{2}, \frac{1}{2}, \frac{3}{8}, \frac{1}{4}, \frac{5}{32}, \dots).$

(4.6) Sei $f_0 := 0$, $f_1 := 1$ und $f_n := f_{n-1} + f_{n-2}$. Dadurch wird rekursiv die Folge der *Fibonacci-Zahlen* definiert:

$$(f_n)_{n \in \mathbb{N}} = (0, 1, 1, 2, 3, 5, 8, 13, 21, 34, \dots).$$

(4.7) Für jede reelle Zahl x hat man die Folge ihrer Potenzen:

$$(x^n)_{n \in \mathbb{N}} = (1, x, x^2, x^3, x^4, \dots).$$

Definition. Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge reeller Zahlen. Die Folge heißt *konvergent* gegen $a \in \mathbb{R}$, falls gilt:

Zu jedem $\varepsilon > 0$ existiert ein $N \in \mathbb{N}$, so dass
 $|a_n - a| < \varepsilon$ für alle $n \geq N$.

Man beachte, dass die Zahl N von ε abhängt. Im Allgemeinen wird man N umso größer wählen müssen, je kleiner ε ist.

Konvergiert (a_n) gegen a , so nennt man a den *Grenzwert* oder den *Limes* der Folge und schreibt

$$\lim_{n \rightarrow \infty} a_n = a \quad \text{oder kurz} \quad \lim a_n = a.$$

Auch die Schreibweise

$$a_n \longrightarrow a \quad \text{für } n \rightarrow \infty$$

(gelesen: a_n strebt gegen a für n gegen unendlich) ist gebräuchlich.

Eine Folge, die gegen 0 konvergiert, nennt man *Nullfolge*.

Geometrische Deutung der Konvergenz. Für $\varepsilon > 0$ versteht man unter der ε -Umgebung von $a \in \mathbb{R}$ die Menge aller Punkte der Zahlengeraden, die von a einen Abstand kleiner als ε haben. Dies ist das Intervall

$$]a - \varepsilon, a + \varepsilon[:= \{x \in \mathbb{R} : a - \varepsilon < x < a + \varepsilon\}.$$

(Die nach außen geöffneten Klammern deuten an, dass die Endpunkte nicht zum Intervall gehören.)

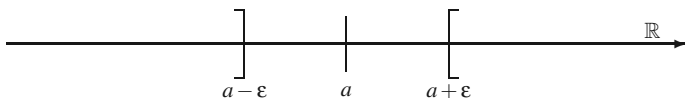


Bild 4.1 ε -Umgebung

Die Konvergenz-Bedingung lässt sich nun so formulieren: Zu jedem $\varepsilon > 0$ existiert ein N , so dass

$$a_n \in]a - \varepsilon, a + \varepsilon[\quad \text{für alle } n \geq N.$$

Die Folge (a_n) konvergiert also genau dann gegen a , wenn in jeder noch so kleinen ε -Umgebung von a fast alle Glieder der Folge liegen. Dabei bedeutet „fast alle“: alle bis auf höchstens endlich viele Ausnahmen.

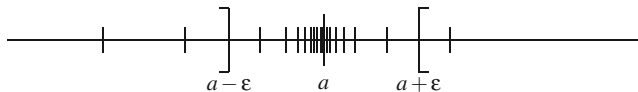


Bild 4.2 Konvergenz

Definition. Eine Folge (a_n) , die nicht konvergiert, heißt *divergent*.

Behandlung der Beispiele

Wir untersuchen jetzt die eingangs gebrachten Beispiele von Folgen auf Konvergenz bzw. Divergenz.

(4.1) Die konstante Folge $(a, a, a, \dots)$ konvergiert trivialerweise gegen a .

(4.2) $\lim_{n \rightarrow \infty} \frac{1}{n} = 0$, die Folge $(1/n)_{n \geq 1}$ ist also eine Nullfolge. Denn sei $\varepsilon > 0$ vorgegeben. Nach dem Archimedischen Axiom gibt es ein $N \in \mathbb{N}$ mit $N > 1/\varepsilon$. Damit ist

$$\left| \frac{1}{n} - 0 \right| = \frac{1}{n} < \varepsilon \quad \text{für alle } n \geq N.$$

Übrigens kann man zeigen, dass die Tatsache, dass $(1/n)_{n \geq 1}$ eine Nullfolge ist, sogar äquivalent mit dem Archimedischen Axiom ist, siehe Aufgabe 4.9.

(4.3) Die Folge $a_n = (-1)^n$, $n \in \mathbb{N}$, divergiert.

Beweis. Angenommen, die Folge (a_n) konvergiert gegen eine reelle Zahl a . Dann gibt es nach Definition zu $\varepsilon := 1$ ein $N \in \mathbb{N}$ mit

$$|a_n - a| < 1 \quad \text{für alle } n \geq N.$$

Für alle $n \geq N$ gilt dann nach der Dreiecks-Ungleichung

$$\begin{aligned} 2 &= |a_{n+1} - a_n| = |(a_{n+1} - a) + (a - a_n)| \\ &\leq |a_{n+1} - a| + |a_n - a| \\ &< 1 + 1 = 2. \end{aligned}$$

Es ergibt sich also der Widerspruch $2 < 2$, d.h. die Folge kann nicht gegen a konvergieren.

(4.4) $\lim_{n \rightarrow \infty} \frac{n}{n+1} = 1$. Zu $\varepsilon > 0$ wählen wir ein $N \in \mathbb{N}$ mit $N > 1/\varepsilon$. Damit ist

$$\left| \frac{n}{n+1} - 1 \right| = \frac{1}{n+1} < \varepsilon \quad \text{für alle } n \geq N.$$

(4.5) $\lim_{n \rightarrow \infty} \frac{n}{2^n} = 0$.

Beweis. Für alle $n > 3$ gilt $n^2 \leq 2^n$, wie man durch vollständige Induktion beweist (vgl. Aufgabe 3.1). Daraus folgt

$$\frac{n^2}{2^n} \leq 1, \quad \text{also} \quad \frac{n}{2^n} \leq \frac{1}{n}.$$

Sei $\varepsilon > 0$ vorgegeben und $N > \max(3, 1/\varepsilon)$. Dann ist

$$\left| \frac{n}{2^n} - 0 \right| = \frac{n}{2^n} \leq \frac{1}{n} < \varepsilon \quad \text{für alle } n \geq N, \quad \text{q.e.d.}$$

Bevor wir die nächsten Beispiele behandeln, führen wir noch einen weiteren wichtigen Begriff ein.

Definition (Beschränktheit von Folgen). Eine Folge $(a_n)_{n \in \mathbb{N}}$ reeller Zahlen heißt *nach oben* (bzw. *nach unten*) *beschränkt*, wenn es eine Konstante $K \in \mathbb{R}$ gibt, so dass

$$a_n \leq K \quad \text{für alle } n \in \mathbb{N} \quad (\text{bzw. } a_n \geq K \quad \text{für alle } n \in \mathbb{N}).$$

Die Folge (a_n) heißt *beschränkt*, wenn es eine reelle Konstante $M \geq 0$ gibt, so dass

$$|a_n| \leq M \quad \text{für alle } n.$$

Bemerkung. Eine Folge (a_n) reeller Zahlen ist genau dann beschränkt, wenn sie sowohl nach oben als auch nach unten beschränkt ist.

Satz 1. Jede konvergente Folge $(a_n)_{n \in \mathbb{N}}$ ist beschränkt.

Beweis. Sei $\lim a_n = a$. Dann gibt es ein $N \in \mathbb{N}$, so dass

$$|a_n - a| < 1 \quad \text{für alle } n \geq N.$$

Daraus folgt

$$|a_n| \leq |a| + |a_n - a| \leq |a| + 1 \quad \text{für } n \geq N.$$

Wir setzen $M := \max(|a_0|, |a_1|, \dots, |a_{N-1}|, |a| + 1)$. Damit gilt

$$|a_n| \leq M \quad \text{für alle } n \in \mathbb{N}, \quad \text{q.e.d.}$$

Bemerkung. Die Umkehrung von Satz 1 gilt nicht. Z.B. ist die Folge $a_n = (-1)^n$, $n \in \mathbb{N}$, beschränkt, aber nicht konvergent.

Wir fahren jetzt mit der Behandlung der Beispiele fort.

(4.6) Die Folge $(f_n) = (0, 1, 1, 2, 3, 5, 8, 13, \dots)$ der Fibonacci-Zahlen divergiert. Denn man zeigt leicht durch vollständige Induktion, dass $f_{n+1} \geq n$ für alle $n \geq 0$. Die Folge ist also nicht beschränkt und kann deshalb nach Satz 1 nicht konvergieren.

Zu den Fibonacci-Zahlen siehe auch Aufgaben 4.2 und 6.6.

(4.7) Das Konvergenzverhalten der Folge $(x^n)_{n \in \mathbb{N}}$ hängt vom Wert von x ab. Wir unterscheiden vier Fälle.

1. Fall. Für $|x| < 1$ gilt $\lim_{n \rightarrow \infty} x^n = 0$.

Beweis. Nach §3, Satz 3 b), existiert zu vorgegebenem $\varepsilon > 0$ ein $N \in \mathbb{N}$ mit $|x|^N < \varepsilon$. Damit ist

$$|x^n - 0| = |x|^n \leq |x|^N < \varepsilon \quad \text{für alle } n \geq N.$$

2. Fall. Für $x = 1$ ist $x^n = 1$ für alle n , also $\lim_{n \rightarrow \infty} x^n = 1$.

3. Fall. $x = -1$. Nach Beispiel (4.3) divergiert die Folge $((-1)^n)_{n \in \mathbb{N}}$.

4. Fall. Für $|x| > 1$ divergiert die Folge (x^n) . Denn aus §3, Satz 3 a), ergibt sich, dass die Folge (x^n) unbeschränkt ist.

Satz 2 (Eindeutigkeit des Limes). *Die Folge (a_n) konvergiere sowohl gegen a als auch gegen b . Dann ist $a = b$.*

Bemerkung. Satz 2 macht die Schreibweise $\lim_{n \rightarrow \infty} a_n = a$ erst sinnvoll.

Beweis. Angenommen, es wäre $a \neq b$. Setze $\varepsilon := |a - b|/2$. Dann gibt es nach Voraussetzung natürliche Zahlen N_1 und N_2 mit

$$|a_n - a| < \varepsilon \quad \text{für } n \geq N_1 \quad \text{und} \quad |a_n - b| < \varepsilon \quad \text{für } n \geq N_2.$$

Für $n := \max(N_1, N_2)$ gilt dann sowohl $|a_n - a| < \varepsilon$ als auch $|a_n - b| < \varepsilon$. Daraus folgt mit der Dreiecks-Ungleichung

$$|a - b| \leq |a - a_n| + |a_n - b| < 2\varepsilon = |a - b|,$$

also der Widerspruch $|a - b| < |a - b|$. Es muss also doch $a = b$ sein.

Häufig benutzt man bei der Untersuchung der Konvergenz von Folgen nicht direkt die Definition, sondern führt die Konvergenz nach gewissen Regeln auf schon bekannte Folgen zurück. Dazu dienen die nächsten Sätze.

Satz 3 (Summe und Produkt konvergenter Folgen). *Seien $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ zwei konvergente Folgen reeller Zahlen. Dann konvergieren auch die Summenfolge $(a_n + b_n)_{n \in \mathbb{N}}$ und die Produktfolge $(a_n b_n)_{n \in \mathbb{N}}$ und es gilt*

$$\lim_{n \rightarrow \infty} (a_n + b_n) = \left(\lim_{n \rightarrow \infty} a_n \right) + \left(\lim_{n \rightarrow \infty} b_n \right),$$

$$\lim_{n \rightarrow \infty} (a_n b_n) = \left(\lim_{n \rightarrow \infty} a_n \right) \cdot \left(\lim_{n \rightarrow \infty} b_n \right).$$

Beweis. Wir bezeichnen die Limites der gegebenen Folgen mit

$$a := \lim_{n \rightarrow \infty} a_n \quad \text{und} \quad b := \lim_{n \rightarrow \infty} b_n.$$

a) Zunächst zur Summenfolge! Es ist zu zeigen

$$\lim_{n \rightarrow \infty} (a_n + b_n) = a + b.$$

Sei $\varepsilon > 0$ vorgegeben. Dann ist auch $\varepsilon/2 > 0$, es gibt also wegen der Konvergenz der Folgen (a_n) und (b_n) Zahlen $N_1, N_2 \in \mathbb{N}$ mit

$$|a_n - a| < \frac{\varepsilon}{2} \quad \text{für } n \geq N_1 \quad \text{und} \quad |b_n - b| < \frac{\varepsilon}{2} \quad \text{für } n \geq N_2.$$

Dann gilt für alle $n \geq N := \max(N_1, N_2)$

$$|(a_n + b_n) - (a + b)| \leq |a_n - a| + |b_n - b| < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon.$$

Damit ist die Konvergenz der Summenfolge bewiesen.

b) Wir zeigen jetzt $\lim_{n \rightarrow \infty} (a_n b_n) = ab$.

Nach Satz 1 ist die Folge (a_n) beschränkt, es gibt also eine reelle Konstante $K > 0$, so dass $|a_n| \leq K$ für alle n . Wir können außerdem (nach evtl. Vergrößerung von K) annehmen, dass $|b| \leq K$. Sei wieder $\varepsilon > 0$ vorgegeben. Da auch $\frac{\varepsilon}{2K} > 0$, gibt es Zahlen $M_1, M_2 \in \mathbb{N}$ mit

$$|a_n - a| < \frac{\varepsilon}{2K} \quad \text{für } n \geq M_1 \quad \text{und} \quad |b_n - b| < \frac{\varepsilon}{2K} \quad \text{für } n \geq M_2.$$

Für alle $n \geq M := \max(M_1, M_2)$ gilt dann

$$|a_n b_n - ab| = |a_n b_n - a_n b + a_n b - ab|$$

$$\begin{aligned}
&= |a_n(b_n - b) + (a_n - a)b| \\
&\leq |a_n||b_n - b| + |a_n - a||b| \\
&< K \cdot \frac{\varepsilon}{2K} + \frac{\varepsilon}{2K} \cdot K = \varepsilon.
\end{aligned}$$

Daraus folgt die Konvergenz der Produktfolge.

Bemerkung. Der hier zur Abschätzung von $|a_nb_n - ab|$ angewandte Trick, einen scheinbar nutzlosen Summanden $0 = -a_nb + a_nb$ einzufügen, wird in der Analysis in ähnlicher Form öfter benutzt.

Corollar (Linearkombination konvergenter Folgen). *Seien $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ zwei konvergente Folgen reeller Zahlen und $\lambda, \mu \in \mathbb{R}$. Dann konvergiert auch die Folge $(\lambda a_n + \mu b_n)_{n \in \mathbb{N}}$ und es gilt*

$$\lim_{n \rightarrow \infty} (\lambda a_n + \mu b_n) = \lambda \lim_{n \rightarrow \infty} a_n + \mu \lim_{n \rightarrow \infty} b_n.$$

Dies ergibt sich aus Satz 3, da man die Folge $(\lambda a_n)_{n \in \mathbb{N}}$ als Produkt der konstanten Folge (λ) mit der Folge (a_n) auffassen kann, und analog für (μb_n) .

Beispielsweise erhält man für $\lambda = 1, \mu = -1$ insbesondere folgende Aussage: Zwei konvergente Folgen (a_n) und (b_n) haben genau dann denselben Grenzwert, wenn die Differenzfolge $(a_n - b_n)$ eine Nullfolge ist.

Satz 4 (Quotient konvergenter Folgen). *Seien $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ zwei konvergente Folgen reeller Zahlen mit $\lim b_n =: b \neq 0$. Dann gibt es ein $n_0 \in \mathbb{N}$, so dass $b_n \neq 0$ für alle $n \geq n_0$ und die Quotientenfolge $(a_n/b_n)_{n \geq n_0}$ konvergiert. Für ihren Grenzwert gilt*

$$\lim_{n \rightarrow \infty} \frac{a_n}{b_n} = \frac{\lim a_n}{\lim b_n}.$$

Beweis. Wir behandeln zunächst den Spezialfall, dass (a_n) die konstante Folge $a_n = 1$ ist. Da $b \neq 0$, ist $|b|/2 > 0$, es gibt also ein $n_0 \in \mathbb{N}$ mit

$$|b_n - b| < \frac{|b|}{2} \quad \text{für alle } n \geq n_0.$$

Daraus folgt $|b_n| \geq |b|/2$, insbesondere $b_n \neq 0$ für $n \geq n_0$. Zu vorgegebenem $\varepsilon > 0$ gibt es ein $N_1 \in \mathbb{N}$, so dass

$$|b_n - b| < \frac{\varepsilon |b|^2}{2} \quad \text{für alle } n \geq N_1.$$

Dann gilt für $n \geq N := \max(n_0, N_1)$

$$\left| \frac{1}{b_n} - \frac{1}{b} \right| = \frac{1}{|b_n||b|} \cdot |b - b_n| < \frac{2}{|b|^2} \cdot \frac{\varepsilon|b|^2}{2} = \varepsilon.$$

Damit ist $\lim(1/b_n) = 1/b$ gezeigt. Der allgemeine Fall folgt mit Satz 3 aus diesem Spezialfall, da sich der Quotient a_n/b_n als Produkt $a_n \cdot (1/b_n)$ schreiben lässt.

(4.8) Wir betrachten als Beispiel die Folge

$$a_n := \frac{3n^2 + 13n}{n^2 - 2}, \quad n \in \mathbb{N}.$$

Für $n > 0$ kann man schreiben $a_n = \frac{3 + 13/n}{1 - 2/n^2}$. Da $\lim(1/n) = 0$, folgt aus Satz 3, dass $\lim(1/n^2) = 0$. Aus dem Corollar zu Satz 3 folgt nun

$$\lim\left(3 + \frac{13}{n}\right) = 3, \quad \text{und} \quad \lim\left(1 - \frac{2}{n^2}\right) = 1.$$

Mit Satz 4 erhält man schließlich

$$\lim_{n \rightarrow \infty} \frac{3n^2 + 13n}{n^2 - 2} = \frac{\lim(3 + \frac{13}{n})}{\lim(1 - \frac{2}{n^2})} = \frac{3}{1} = 3.$$

Satz 5 (Größenvergleich konvergenter Folgen). *Seien (a_n) und (b_n) zwei konvergente Folgen reeller Zahlen mit $a_n \leq b_n$ für alle n . Dann gilt auch*

$$\lim_{n \rightarrow \infty} a_n \leq \lim_{n \rightarrow \infty} b_n.$$

Vorsicht! Wenn $a_n < b_n$ für alle n , dann ist nicht notwendig $\lim a_n < \lim b_n$, wie man an dem Beispiel der Folgen $a_n = 0$ und $b_n = \frac{1}{n}$, ($n \geq 1$), sieht, die beide gegen 0 konvergieren.

Beweis. Durch Übergang zur Differenzenfolge $(b_n - a_n)$ genügt es nach dem Corollar zu Satz 3, folgendes zu beweisen: Ist (c_n) eine konvergente Folge mit $c_n \geq 0$ für alle n , so gilt auch $\lim c_n \geq 0$.

Hierfür geben wir einen Widerspruchsbeweis. Wäre dies nicht der Fall, so hätten wir

$$\lim_{n \rightarrow \infty} c_n = -\varepsilon \quad \text{mit einem } \varepsilon > 0$$

und es gäbe ein $N \in \mathbb{N}$ mit $|c_n - (-\varepsilon)| < \varepsilon$ für alle $n \geq N$, woraus der Widerspruch $c_n < 0$ für $n \geq N$ folgen würde.

Corollar. Seien $A \leq B$ reelle Zahlen und (a_n) eine konvergente Folge mit $A \leq a_n \leq B$ für alle n . Dann gilt auch

$$A \leq \lim_{n \rightarrow \infty} a_n \leq B.$$

Unendliche Reihen

Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge reeller Zahlen. Daraus entsteht eine (unendliche) Reihe, indem man, grob gesprochen, die Folgenglieder durch ein Pluszeichen verbindet:

$$a_0 + a_1 + a_2 + a_3 + a_4 + \dots$$

Dies lässt sich so präzisieren: Für jedes $m \in \mathbb{N}$ betrachte man die sog. *Partialsumme*

$$s_m := \sum_{n=0}^m a_n = a_0 + a_1 + a_2 + \dots + a_m.$$

Die Folge $(s_m)_{m \in \mathbb{N}}$ der Partialsummen heißt (unendliche) *Reihe* mit den Gliedern a_n und wird mit $\sum_{n=0}^{\infty} a_n$ bezeichnet. Konvergiert die Folge $(s_m)_{m \in \mathbb{N}}$ der Partialsummen, so wird ihr Grenzwert ebenfalls mit $\sum_{n=0}^{\infty} a_n$ bezeichnet und heißt dann *Summe der Reihe*.

Das Symbol $\sum_{n=0}^{\infty} a_n$ bedeutet also zweierlei:

- i) Die Folge $\left(\sum_{n=0}^m a_n \right)_{m \in \mathbb{N}}$ der Partialsummen.
- ii) Im Falle der Konvergenz den Grenzwert $\lim_{m \rightarrow \infty} \sum_{n=0}^m a_n$.

Entsprechend sind natürlich Reihen $\sum_{n=k}^{\infty} a_n$ definiert, bei denen die Indexmenge nicht bei 0 beginnt.

Übrigens lässt sich jede Folge $(c_n)_{n \in \mathbb{N}}$ auch als Reihe darstellen, denn es gilt

$$c_n = c_0 + \sum_{k=1}^n (c_k - c_{k-1}) \quad \text{für alle } n \in \mathbb{N}.$$

Eine solche Darstellung, in der sich zwei aufeinander folgende Terme immer zur Hälfte wegekürzen, nennt man auch *Teleskop-Summe*.

(4.9) Beispiel. Mit $c_n := \frac{n}{n+1}$ ist $c_0 = 0$ und

$$c_k - c_{k-1} = \frac{k}{k+1} - \frac{k-1}{k} = \frac{1}{k(k+1)}.$$

Deshalb gilt $\sum_{k=1}^n \frac{1}{k(k+1)} = \frac{n}{n+1}$ und

$$\sum_{k=1}^{\infty} \frac{1}{k(k+1)} = \lim_{n \rightarrow \infty} \frac{n}{n+1} = 1.$$

Satz 6 (Unendliche geometrische Reihe). *Die Reihe $\sum_{n=0}^{\infty} x^n$ konvergiert für alle $|x| < 1$ mit dem Grenzwert*

$$\sum_{n=0}^{\infty} x^n = \frac{1}{1-x}.$$

Beweis. Für die Partialsummen gilt nach §1, Satz 6

$$s_n = \sum_{k=0}^n x^k = \frac{1-x^{n+1}}{1-x}.$$

Nach Beispiel (4.7) ist $\lim_{n \rightarrow \infty} x^{n+1} = 0$, also $\lim s_n = \frac{1}{1-x}$, q.e.d.

(4.10) Beispiele. Für $x = \pm \frac{1}{2}$ erhält man die beiden Formeln

$$1 + \frac{1}{2} + \frac{1}{4} + \frac{1}{8} + \frac{1}{16} + \dots = \frac{1}{1-1/2} = 2,$$

$$1 - \frac{1}{2} + \frac{1}{4} - \frac{1}{8} + \frac{1}{16} \mp \dots = \frac{1}{1+1/2} = \frac{2}{3}.$$

Satz 7 (Linearkombination konvergenter Reihen). *Seien*

$$\sum_{n=0}^{\infty} a_n \quad \text{und} \quad \sum_{n=0}^{\infty} b_n$$

zwei konvergente Reihen reeller Zahlen und $\lambda, \mu \in \mathbb{R}$. Dann konvergiert auch die Reihe $\sum_{n=0}^{\infty} (\lambda a_n + \mu b_n)$ und es gilt

$$\sum_{n=0}^{\infty} (\lambda a_n + \mu b_n) = \lambda \sum_{n=0}^{\infty} a_n + \mu \sum_{n=0}^{\infty} b_n.$$

Dies ergibt sich sofort, wenn man das Corollar zu Satz 3 auf die Partialsummen anwendet.

Bemerkung. Mit den Begriffen aus der Linearen Algebra lässt sich Satz 7 abstrakt so interpretieren: Die konvergenten Reihen bilden einen Vektorraum über dem Körper $\mathbb{R}$, und die Abbildung, die einer konvergenten Reihe ihre Summe zuordnet, ist eine Linearform auf diesem Vektorraum.

Bei konvergenten Folgen hatten wir auch eine einfache Aussage über Produkte. Im Gegensatz dazu sind die Verhältnisse bei Produkten konvergenter Reihen viel komplizierter. Wir werden uns damit in §8 beschäftigen.

(4.11) Unendliche Dezimalbrüche sind spezielle Reihen. Wir betrachten hier als Beispiel den *periodischen Dezimalbruch*

$$x := 0.086363\overline{63},$$

wobei die Überstreichung von 63 andeuten soll, dass sich diese Zifferngruppe unendlich oft wiederholt. Dies bedeutet, dass x den folgenden Wert hat:

$$x = \frac{8}{100} + \frac{63}{10^4} + \frac{63}{10^6} + \dots = \frac{8}{100} + \sum_{k=0}^{\infty} \frac{63}{10^{4+2k}}.$$

Nach den Sätzen 6 und 7 ist

$$\sum_{k=0}^{\infty} \frac{63}{10^{4+2k}} = \frac{63}{10^4} \sum_{k=0}^{\infty} (10^{-2})^k = \frac{63}{10^4} \cdot \frac{1}{1 - 10^{-2}} = \frac{63}{9900},$$

also

$$x = \frac{8}{100} + \frac{63}{9900} = \frac{855}{9900} = \frac{19}{220}.$$

Im nächsten Paragraphen werden wir uns systematischer mit unendlichen Dezimalbrüchen beschäftigen.

Bestimmte Divergenz gegen $\pm\infty$

Definition. Eine Folge $(a_n)_{n \in \mathbb{N}}$ reeller Zahlen heißt *bestimmt divergent* gegen $+\infty$, wenn zu jedem $K \in \mathbb{R}$ ein $N \in \mathbb{N}$ existiert, so dass

$$a_n > K \quad \text{für alle } n \geq N.$$

Die Folge (a_n) heißt *bestimmt divergent* gegen $-\infty$, wenn die Folge $(-a_n)$ bestimmt gegen $+\infty$ divergiert.

Divergiert (a_n) bestimmt gegen $+\infty$ (bzw. $-\infty$), so schreibt man

$$\lim_{n \rightarrow \infty} a_n = \infty, \quad (\text{bzw. } \lim_{n \rightarrow \infty} a_n = -\infty).$$

Statt bestimmt divergent sagt man auch *uneigentlich konvergent*.

Beispiele

(4.12) Die Folge $a_n = n$, $n \in \mathbb{N}$, divergiert bestimmt gegen $+\infty$.

(4.13) Die Folge $a_n = -2^n$, $n \in \mathbb{N}$, divergiert bestimmt gegen $-\infty$.

(4.14) Die Folge $a_n = (-1)^n n$, $n \in \mathbb{N}$, divergiert. Sie divergiert jedoch weder bestimmt gegen $+\infty$ noch bestimmt gegen $-\infty$.

Bemerkungen. a) Wie aus der Definition unmittelbar folgt, ist eine Folge, die bestimmt gegen $+\infty$ (bzw. $-\infty$) divergiert, nicht nach oben (bzw. nicht nach unten) beschränkt. Die Umkehrung gilt jedoch nicht, wie Beispiel (4.14) zeigt.

b) $+\infty$ und $-\infty$ sind Symbole, deren Bedeutung durch die Definition der bestimmten Divergenz genau festgelegt ist. Sie lassen sich nicht als reelle Zahlen auffassen, sonst ergäben sich Widersprüche. Sei etwa $a_n := n$, $b_n := 1$ und $c_n := a_n + b_n = n + 1$. Dann ist $\lim a_n = \infty$, $\lim b_n = 1$ und $\lim c_n = \infty$. Könnte man mit ∞ so rechnen wie mit reellen Zahlen, würde nach Satz 3 gelten $\infty + 1 = \infty$. Nach (2.4) besitzt die Gleichung $a + x = a$ die eindeutige Lösung $x = 0$. Man erhielte damit den Widerspruch $1 = 0$.

Es ist jedoch für manche Zwecke nützlich, die sog. erweiterte Zahlengerade $\overline{\mathbb{R}} := \mathbb{R} \cup \{+\infty, -\infty\}$ einzuführen und

$$-\infty < x < +\infty \quad \text{für alle } x \in \mathbb{R}$$

zu definieren.

Die nächsten beiden Sätze stellen eine Beziehung zwischen der bestimmten Divergenz gegen $\pm\infty$ und der Konvergenz gegen 0 her.

Satz 8 (Reziprokes einer bestimmt divergenten Folge). *Die Folge $(a_n)_{n \in \mathbb{N}}$ sei bestimmt divergent gegen $+\infty$ oder $-\infty$. Dann gibt es ein $n_0 \in \mathbb{N}$, so dass $a_n \neq 0$ für alle $n \geq n_0$, und es gilt*

$$\lim_{n \rightarrow \infty} \frac{1}{a_n} = 0.$$

Beweis. Sei $\lim a_n = +\infty$. Dann gibt es nach Definition zur Schranke $K = 0$ ein $n_0 \in \mathbb{N}$ mit $a_n > 0$ für alle $n \geq n_0$. Insbesondere ist $a_n \neq 0$ für $n \geq n_0$.

Wir zeigen jetzt $\lim(1/a_n) = 0$. Sei $\varepsilon > 0$ vorgegeben. Da $\lim a_n = \infty$, gibt es ein $N \in \mathbb{N}$ mit $a_n > 1/\varepsilon$ für alle $n \geq N$. Daraus folgt $1/a_n < \varepsilon$ für alle $n \geq N$, q.e.d.

Der Fall $\lim a_n = -\infty$ wird durch Übergang zur Folge $(-a_n)$ bewiesen.

Satz 9 (Reziprokes einer Nullfolge). *Sei $(a_n)_{n \in \mathbb{N}}$ eine Nullfolge mit $a_n > 0$ für alle n (bzw. $a_n < 0$ für alle n). Dann divergiert die Folge $(1/a_n)_{n \in \mathbb{N}}$ bestimmt gegen $+\infty$ (bzw. gegen $-\infty$).*

Beweis. Wir behandeln nur den Fall einer positiven Nullfolge. Sei $K > 0$ eine vorgegebene Schranke. Wegen $\lim a_n = 0$ gibt es ein $N \in \mathbb{N}$, so dass

$$|a_n| < \varepsilon := \frac{1}{K} \quad \text{für alle } n \geq N.$$

Also ist $1/a_n = 1/|a_n| > K$ für alle $n \geq N$, d.h. $\lim(1/a_n) = \infty$.

(4.15) Beispielsweise ist $\lim_{n \rightarrow \infty} \frac{2^n}{n} = \infty$, wie aus (4.5) folgt.

AUFGABEN

4.1. Seien a und b reelle Zahlen. Die Folge $(a_n)_{n \in \mathbb{N}}$ sei wie folgt rekursiv definiert:

$$a_0 := a, \quad a_1 := b, \quad a_n := \frac{1}{2}(a_{n-1} + a_{n-2}) \quad \text{für } n \geq 2.$$

Man beweise, dass die Folge $(a_n)_{n \in \mathbb{N}}$ konvergiert und bestimme ihren Grenzwert.

4.2. a) Für die in (4.6) definierten Fibonacci-Zahlen beweise man

$$f_{n+1}f_{n-1} - f_n^2 = (-1)^n \quad \text{für alle } n \geq 1.$$

b) Man zeige $\lim_{n \rightarrow \infty} \frac{f_{n+1}f_{n-1}}{f_n^2} = 1$.

4.3. Man berechne die Summe der Reihe

$$\sum_{n=1}^{\infty} \frac{1}{4n^2 - 1}.$$

4.4. Man berechne das unendliche Produkt

$$\prod_{n=2}^{\infty} \frac{n^3 - 1}{n^3 + 1},$$

d.h. den Limes der Folge $p_m := \prod_{n=2}^m \frac{n^3 - 1}{n^3 + 1}$, $m \geq 2$.

4.5. a) Es sei $(a_n)_{n \in \mathbb{N}}$ eine Folge, die gegen ein $a \in \mathbb{R}$ konvergiere. Man beweise, dass dann die Folge $(b_n)_{n \in \mathbb{N}}$, definiert durch

$$b_n := \frac{1}{n+1}(a_0 + a_1 + \dots + a_n) \quad \text{für alle } n \in \mathbb{N},$$

ebenfalls gegen a konvergiert.

b) Man gebe ein Beispiel einer nicht konvergenten Folge $(a_n)_{n \in \mathbb{N}}$ an, bei dem die wie in a) definierte Folge (b_n) konvergiert.

4.6. Man beweise: Für jede reelle Zahl $b > 1$ und jede natürliche Zahl k gilt

$$\lim_{n \rightarrow \infty} \frac{b^n}{n^k} = \infty.$$

4.7. Seien $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ Folgen reeller Zahlen mit $\lim a_n = \infty$ und $\lim b_n =: b \in \mathbb{R}$. Man beweise:

a) $\lim (a_n + b_n) = \infty$.

b) Ist $b > 0$, so gilt $\lim (a_n b_n) = \infty$; ist $b < 0$, so gilt $\lim (a_n b_n) = -\infty$.

4.8. Man gebe Beispiele reeller Zahlenfolgen $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ mit $\lim a_n = \infty$ und $\lim b_n = 0$ an, so dass jeder der folgenden Fälle eintritt:

a) $\lim (a_n b_n) = +\infty$.

b) $\lim (a_n b_n) = -\infty$.

c) $\lim (a_n b_n) = c$, wobei c eine beliebig vorgegebene reelle Zahl ist.

d) Die Folge $(a_n b_n)_{n \in \mathbb{N}}$ ist beschränkt, aber nicht konvergent.

4.9. Man beweise: In einem angeordneten Körper ist jede der folgenden Aussagen äquivalent zum Archimedischen Axiom:

(1) Die Folge $(1/n)_{n \geq 1}$ ist eine Nullfolge.

(2) Die Folge $(2^{-n})_{n \in \mathbb{N}}$ ist eine Nullfolge.

§ 5 Das Vollständigkeits-Axiom

Mithilfe der bisher behandelten Axiome lässt sich nicht die Existenz von Irrationalzahlen beweisen, denn all diese Axiome gelten auch im Körper der rationalen Zahlen. Bekanntlich gibt es (was schon die alten Griechen wussten) keine rationale Zahl, deren Quadrat gleich 2 ist. Also lässt sich mit den bisherigen Axiomen nicht beweisen, dass eine Quadratwurzel aus 2 existiert. Es ist ein weiteres Axiom nötig, das sogenannte Vollständigkeits-Axiom. Aus diesem folgt unter anderem, dass jeder unendliche Dezimalbruch (ob periodisch oder nicht) gegen eine reelle Zahl konvergiert.

Eine charakteristische Eigenschaft konvergenter Folgen, die formuliert werden kann, ohne auf den Grenzwert der Folge Bezug zu nehmen, wurde von Cauchy entdeckt.

Definition. Eine Folge $(a_n)_{n \in \mathbb{N}}$ reeller Zahlen heißt *Cauchy-Folge*, wenn gilt: Zu jedem $\varepsilon > 0$ existiert ein $N \in \mathbb{N}$, so dass

$$|a_n - a_m| < \varepsilon \quad \text{für alle } n, m \geq N.$$

Eine andere Bezeichnung für Cauchy-Folge ist *Fundamental-Folge*.

Grob gesprochen kann man also sagen: Eine Folge ist eine Cauchy-Folge, wenn die Folgenglieder untereinander beliebig wenig abweichen, falls nur die Indizes genügend groß sind. Man beachte: Es genügt nicht, dass die Differenz $|a_n - a_{n+1}|$ zweier aufeinander folgender Folgenglieder beliebig klein wird, sondern die Differenz $|a_n - a_m|$ muss kleiner als ein beliebiges $\varepsilon > 0$ sein, wobei n und m unabhängig voneinander alle natürlichen Zahlen durchlaufen, die größer-gleich einer von ε abhängigen Schranke sind. Bei konvergenten Folgen ist das der Fall, wie der nächste Satz zeigt.

Satz 1. *Jede konvergente Folge reeller Zahlen ist eine Cauchy-Folge.*

Beweis. Die Folge (a_n) konvergiere gegen a . Dann gibt es zu vorgegebenem $\varepsilon > 0$ ein $N \in \mathbb{N}$, so dass

$$|a_n - a| < \frac{\varepsilon}{2} \quad \text{für alle } n \geq N.$$

Für alle $n, m \geq N$ gilt dann

$$\begin{aligned}
 |a_n - a_m| &= |(a_n - a) - (a_m - a)| \\
 &\leq |a_n - a| + |a_m - a| < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon, \quad \text{q.e.d.}
 \end{aligned}$$

Die Umkehrung von Satz 1 formulieren wir nun als Axiom.

Vollständigkeits-Axiom. *In $\mathbb{R}$ konvergiert jede Cauchy-Folge.*

Bemerkung. Wir werden im nächsten Paragraphen mithilfe des Vollständigkeits-Axioms die Existenz der Quadratwurzeln aus jeder positiven reellen Zahl beweisen. Dies ist mit den bisherigen Axiomen allein noch nicht möglich. Denn da diese auch im Körper der rationalen Zahlen gelten, würde dann z.B. folgen, dass die Quadratwurzel aus 2 rational ist, was aber falsch ist. Also ist das Vollständigkeits-Axiom unabhängig von den bisherigen Axiomen.

Wir erinnern kurz an den wohl aus der Schule bekannten Beweis der Irrationalität der Quadratwurzel aus 2. Wäre diese rational, gäbe es ganze Zahlen $n, m > 0$ mit $(n/m)^2 = 2$. Wir nehmen den Bruch n/m in gekürzter Form an und können deshalb voraussetzen, dass höchstens eine der beiden Zahlen n, m gerade ist. Aus der obigen Gleichung folgt $n^2 = 2m^2$, also ist n gerade, d.h. $n = 2k$ mit einer ganzen Zahl k . Einsetzen und Kürzen ergibt $2k^2 = m^2$, woraus folgt, dass auch m gerade sein muss, Widerspruch!

Das Vollständigkeits-Axiom ist nicht besonders anschaulich. Wir wollen deshalb zeigen, dass es zu einer sehr anschaulichen Aussage, nämlich dem Intervallschachtelungs-Prinzip, äquivalent ist. Sind $a \leq b$ reelle Zahlen, so versteht man unter dem abgeschlossenen Intervall mit Endpunkten a und b die Menge aller Punkte auf der reellen Zahlengeraden, die zwischen a und b liegen, wobei die Endpunkte mit eingeschlossen seien:

$$[a, b] := \{x \in \mathbb{R} : a \leq x \leq b\}.$$

Die Länge (oder der Durchmesser) des Intervalls wird durch

$$\text{diam}([a, b]) := b - a$$

definiert. Damit können wir formulieren:

Intervallschachtelungs-Prinzip. *Sei*

$$I_0 \supset I_1 \supset I_2 \supset \dots \supset I_n \supset I_{n+1} \supset \dots$$

eine absteigende Folge von abgeschlossenen Intervallen in $\mathbb{R}$ mit

$$\lim_{n \rightarrow \infty} \text{diam}(I_n) = 0.$$

Dann gibt es genau eine reelle Zahl x mit $x \in I_n$ für alle $n \in \mathbb{N}$.

Man hat sich vorzustellen, dass die ineinander geschachtelten Intervalle auf den Punkt x „zusammenschrumpfen“, siehe Bild 5.1.

Wir zeigen nun in zwei Schritten die Gleichwertigkeit des Vollständigkeits-Axioms mit dem Intervallschachtelungs-Prinzip.

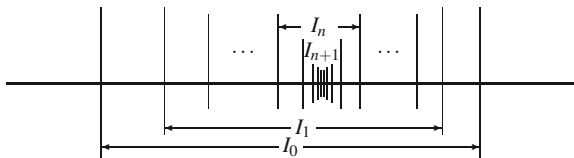


Bild 5.1 Intervallschachtelung

Satz 2. Das Vollständigkeits-Axiom impliziert das Intervallschachtelungs-Prinzip.

Beweis. Seien $I_n = [a_n, b_n]$, $n \in \mathbb{N}$, die ineinander geschachtelten Intervalle. Wir zeigen zunächst, dass die Folge (a_n) der linken Endpunkte eine Cauchy-Folge darstellt.

Beweis hierfür. Da die Länge der Intervalle gegen null konvergiert, gibt es zu vorgegebenem $\varepsilon > 0$ ein $N \in \mathbb{N}$, so dass

$$\text{diam}(I_n) < \varepsilon \quad \text{für alle } n \geq N.$$

Sind $n, m \geq N$, so liegen die Punkte a_n und a_m beide im Intervall I_N , woraus folgt

$$|a_n - a_m| \leq \text{diam}(I_N) < \varepsilon, \quad \text{q.e.d.}$$

Nach dem Vollständigkeits-Axiom konvergiert die Folge (a_n) gegen einen Punkt $x \in \mathbb{R}$. Da $a_k \leq a_n \leq b_n \leq b_k$ für alle $n \geq k$, folgt aus §4, Corollar zu Satz 5, dass $a_k \leq x \leq b_k$. Das heißt, dass der Grenzwert x in allen Intervallen I_k enthalten ist. Da die Länge der Intervalle gegen null konvergiert, kann es nicht mehr als einen solchen Punkt geben. Damit ist Satz 2 bewiesen.

Satz 3. Das Intervallschachtelungs-Prinzip impliziert das Vollständigkeits-Axiom.

Beweis. Sei $(a_n)_{n \in \mathbb{N}}$ eine vorgegebene Cauchy-Folge. Nach Definition gibt es eine Folge $n_0 < n_1 < n_2 < \dots$ natürlicher Zahlen mit

$$|a_n - a_m| < 2^{-k} \quad \text{für alle } n, m \geq n_k.$$

Wir definieren nun

$$I_k := \{x \in \mathbb{R} : |x - a_{n_k}| \leq 2^{-k+1}\}.$$

Die I_k sind abgeschlossene Intervalle mit $I_k \supset I_{k+1}$ für alle k . Denn sei etwa $x \in I_{k+1}$. Dann ist $|x - a_{n_{k+1}}| \leq 2^{-k}$; außerdem ist

$$|a_{n_{k+1}} - a_{n_k}| < 2^{-k},$$

woraus nach der Dreiecks-Ungleichung folgt $|x - a_{n_k}| < 2^{-k+1}$, d.h. $x \in I_k$. Da die Längen der Intervalle gegen null konvergieren, können wir das Intervallschachtelungs-Prinzip anwenden und erhalten einen Punkt $x_0 \in \mathbb{R}$, der in allen I_k liegt, d.h.

$$|x_0 - a_{n_k}| \leq 2^{-k+1} \quad \text{für alle } k \geq 0.$$

Für $n \geq n_k$ ist $|a_n - a_{n_k}| < 2^{-k}$, also insgesamt

$$|x_0 - a_n| < 2^{-k+1} + 2^{-k} < 2^{-k+2},$$

woraus folgt $\lim_{n \rightarrow \infty} a_n = x_0$, die Cauchy-Folge konvergiert also. Damit ist Satz 3 bewiesen.

Wegen der bewiesenen Äquivalenz hätten wir statt des Axioms über die Konvergenz von Cauchy-Folgen auch das Intervallschachtelungs-Prinzip zum Axiom erheben können. Wir haben das Vollständigkeits-Axiom mit den Cauchy-Folgen gewählt, da diese einen zentralen Begriff in der Analysis darstellen, der auch noch in viel allgemeineren Situationen anwendbar ist. (So wird die Leseerin, die tiefer in das Studium der Analysis einsteigt, später sicherlich auf den Begriff des vollständigen metrischen Raumes und des vollständigen topologischen Vektorraums stoßen. In beiden Fällen wird die Vollständigkeit mithilfe von Cauchy-Folgen definiert.)

***b*-adische Brüche**

Sei b eine natürliche Zahl ≥ 2 . Unter einem (unendlichen) b -adischen Bruch versteht man eine Reihe der Gestalt

$$\pm \sum_{n=-k}^{\infty} a_n b^{-n}.$$

Dabei ist $k \geq 0$ und die a_n sind natürliche Zahlen mit $0 \leq a_n < b$. Falls die Basis festgelegt ist, kann man einen b -adischen Bruch auch einfach durch die Aneinanderreihung der Ziffern a_n angeben:

$$\pm a_{-k} a_{-k+1} \cdots a_{-1} a_0 . a_1 a_2 a_3 a_4 a_5 \cdots$$

Dabei werden die Koeffizienten der negativen Potenzen der Basis b durch einen Punkt von den Koeffizienten der nicht-negativen Potenzen abgetrennt. Falls von einer Stelle $k_0 \geq 1$ an alle Koeffizienten $a_k = 0$ sind, lässt man diese auch weg und erhält einen endlichen b -adischen Bruch.

Für $b = 10$ spricht man von Dezimalbrüchen. Im Fall $b = 2$ (dyadische Brüche) sind nur die Ziffern 0 und 1 nötig. Dies eignet sich besonders gut für die interne Darstellung von Zahlen im Computer. Die Babylonier haben das Sexagesimalsystem ($b = 60$) verwendet.

Satz 4. *Jeder b -adische Bruch stellt eine Cauchy-Folge dar, konvergiert also gegen eine reelle Zahl.*

Beweis. Es genügt, einen nicht-negativen b -adischen Bruch $\sum_{n=-k}^{\infty} a_n b^{-n}$ zu betrachten. Für $n \geq -k$ bezeichnen wir die Partialsummen mit

$$x_n := \sum_{v=-k}^n a_v b^{-v}.$$

Wir haben zu zeigen, dass $(x_n)_{n \geq -k}$ eine Cauchy-Folge ist. Sei $\varepsilon > 0$ vorgegeben und $N \in \mathbb{N}$ so groß, dass $b^{-N} < \varepsilon$. Dann gilt für $n \geq m \geq N$

$$\begin{aligned} |x_n - x_m| &= \sum_{v=m+1}^n a_v b^{-v} \leq \sum_{v=m+1}^n (b-1) b^{-v} \\ &\leq (b-1) b^{-m-1} \sum_{v=0}^{n-m-1} b^{-v} \\ &< (b-1) b^{-m-1} \frac{1}{1-b^{-1}} = b^{-m} \leq b^{-N} < \varepsilon. \end{aligned}$$

Damit ist die Behauptung bewiesen.

Von Satz 4 gilt auch die Umkehrung.

Satz 5. *Sei b eine natürliche Zahl ≥ 2 . Dann lässt sich jede reelle Zahl in einen b -adischen Bruch entwickeln.*

Bemerkung. Aus Satz 5 folgt insbesondere, dass sich jede reelle Zahl beliebig genau durch rationale Zahlen approximieren lässt, denn die Partialsummen eines b -adischen Bruches sind rational.

Beweis. Es genügt, den Satz für reelle Zahlen $x \geq 0$ zu beweisen. Nach §3, Satz 3, gibt es mindestens eine natürliche Zahl m mit $x < b^{m+1}$. Sei k die kleinste

natürliche Zahl, so dass

$$0 \leq x < b^{k+1}.$$

Wir konstruieren jetzt durch vollständige Induktion eine Folge $(a_v)_{v \geq -k}$ natürlicher Zahlen $0 \leq a_v < b$, so dass für alle $n \geq -k$ gilt

$$x = \sum_{v=-k}^n a_v b^{-v} + \xi_n \quad \text{mit } 0 \leq \xi_n < b^{-n}.$$

Wegen $\lim_{n \rightarrow \infty} \xi_n = 0$ folgt dann $x = \sum_{v=-k}^{\infty} a_v b^{-v}$, also die Behauptung.

Induktionsanfang $n = -k$. Es gilt $0 \leq x b^{-k} < b$, also gibt es eine ganze Zahl $a_{-k} \in \{0, 1, \dots, b-1\}$ und eine reelle Zahl δ mit $0 \leq \delta < 1$, so dass $x b^{-k} = a_{-k} + \delta$. Mit $\xi_{-k} := \delta b^k$ erhält man

$$x = a_{-k} b^k + \xi_{-k} \quad \text{mit } 0 \leq \xi_{-k} < b^k.$$

Das ist die Behauptung für $n = -k$.

Induktionsschritt $n \rightarrow n+1$. Es gilt $0 \leq \xi_n b^{n+1} < b$, also gibt es eine ganze Zahl $a_{n+1} \in \{0, 1, \dots, b-1\}$ und eine reelle Zahl δ mit $0 \leq \delta < 1$, so dass $\xi_n b^{n+1} = a_{n+1} + \delta$. Mit $\xi_{n+1} := \delta b^{-n-1}$ erhält man

$$x = \sum_{v=-k}^n a_v b^{-v} + (a_{n+1} + \delta) b^{-n-1} = \sum_{v=-k}^{n+1} a_v b^{-v} + \xi_{n+1},$$

wobei $0 \leq \xi_{n+1} < b^{-n-1}$, q.e.d.

Bemerkung. Die Sätze 4 und 5 sagen insbesondere, dass sich jede reelle Zahl durch einen (unendlichen) Dezimalbruch darstellen lässt und umgekehrt. Wir haben also, ausgehend von den Axiomen, die gewohnte Darstellung der reellen Zahlen wiedergefunden.

Man beachte, dass die Darstellung einer reellen Zahl durch einen b -adischen Bruch nicht immer eindeutig ist. Beispielsweise stellen die Dezimalbrüche $1.000000\dots$ und $0.999999\dots$ beide die Zahl 1 dar, denn nach der Summenformel für die unendliche geometrische Reihe ist

$$\sum_{k=1}^{\infty} 9 \cdot 10^{-k} = \frac{9}{10} \sum_{k=0}^{\infty} \left(\frac{1}{10}\right)^k = \frac{9}{10} \cdot \frac{1}{1 - 1/10} = 1.$$

Das hier gegebene Beispiel für die Mehrdeutigkeit ist typisch für den allgemeinen Fall, siehe Aufgabe 5.3.

Teilfolgen

Definition. Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge und

$$n_0 < n_1 < n_2 < \dots$$

eine aufsteigende Folge natürlicher Zahlen. Dann heißt die Folge

$$(a_{n_k})_{k \in \mathbb{N}} = (a_{n_0}, a_{n_1}, a_{n_2}, \dots)$$

Teilfolge der Folge (a_n) .

Es folgt unmittelbar aus der Definition: Ist $(a_n)_{n \in \mathbb{N}}$ eine konvergente Folge mit dem Limes a , so konvergiert auch jede Teilfolge gegen a . Schwieriger ist das Problem, aus nicht-konvergenten Folgen konvergente Teilfolgen zu konstruieren. Die wichtigste Aussage in dieser Richtung ist der folgende Satz.

Satz 6 (Bolzano-Weierstraß). *Jede beschränkte Folge $(a_n)_{n \in \mathbb{N}}$ reeller Zahlen besitzt eine konvergente Teilfolge.*

Beweis. a) Da die Folge beschränkt ist, gibt es Zahlen $A, B \in \mathbb{R}$ mit $A \leq a_n \leq B$ für alle $n \in \mathbb{N}$. Die ganze Folge ist also in dem Intervall

$$[A, B] := \{x \in \mathbb{R} : A \leq x \leq B\}$$

enthalten. Wir konstruieren nun durch vollständige Induktion eine Folge von abgeschlossenen Intervallen $I_k \subset \mathbb{R}$, $k \in \mathbb{N}$, mit folgenden Eigenschaften:

- i) In I_k liegen unendlich viele Glieder der Folge (a_n) ,
- ii) $I_k \subset I_{k-1}$ für $k \geq 1$,
- iii) $\text{diam}(I_k) = 2^{-k} \text{diam}(I_0)$.

Für den *Induktionsanfang* können wir das Intervall $I_0 := [A, B]$ wählen.

Induktionsschritt $k \rightarrow k+1$. Sei das Intervall $I_k = [A_k, B_k]$ mit den Eigenschaften i) bis iii) bereits konstruiert. Sei $M := (A_k + B_k)/2$ die Mitte des Intervalls. Da in I_k unendlich viele Glieder der Folge liegen, muss mindestens eines der Teilintervalle $[A_k, M]$ und $[M, B_k]$ unendlich viele Folgenglieder enthalten. Wir setzen $I_{k+1} := [A_k, M]$, falls in diesem Intervall unendlich viele Folgenglieder liegen, sonst $I_{k+1} := [M, B_k]$. Offenbar hat I_{k+1} wieder die Eigenschaften i) bis iii).

b) Wir definieren nun induktiv eine Teilfolge $(a_{n_k})_{k \in \mathbb{N}}$ mit $a_{n_k} \in I_k$ für alle $k \in \mathbb{N}$.

Induktionsanfang. Wir setzen $n_0 := 0$, d.h. $a_{n_0} = a_0$.

Induktionsschritt $k \rightarrow k+1$. Da in dem Intervall I_{k+1} unendlich viele Glieder der Folge (a_n) liegen, gibt es ein $n_{k+1} > n_k$ mit $a_{n_{k+1}} \in I_{k+1}$.

c) Wir beweisen nun, dass die Teilfolge (a_{n_k}) konvergiert, indem wir zeigen, dass sie eine Cauchy-Folge ist.

Sei $\varepsilon > 0$ vorgegeben und N so groß gewählt, dass $\text{diam}(I_N) < \varepsilon$. Dann gilt für alle $k, j \geq N$

$$a_{n_k} \in I_k \subset I_N \quad \text{und} \quad a_{n_j} \in I_j \subset I_N.$$

Also ist

$$|a_{n_k} - a_{n_j}| \leq \text{diam}(I_N) < \varepsilon, \quad \text{q.e.d.}$$

Definition. Eine Zahl a heißt *Häufungspunkt* einer Folge $(a_n)_{n \in \mathbb{N}}$, wenn es eine Teilfolge von (a_n) gibt, die gegen a konvergiert.

Mit dieser Definition kann man den Inhalt des Satzes von Bolzano-Weierstraß auch so ausdrücken: Jede beschränkte Folge reeller Zahlen besitzt mindestens einen Häufungspunkt.

Wir geben einige Beispiele für Häufungspunkte.

(5.1) Die durch $a_n := (-1)^n$ definierte Folge (a_n) besitzt die Häufungspunkte $+1$ und -1 . Denn es gilt

$$\lim_{k \rightarrow \infty} a_{2k} = 1 \quad \text{und} \quad \lim_{k \rightarrow \infty} a_{2k+1} = -1.$$

(5.2) Die Folge $a_n := (-1)^n + \frac{1}{n}$, $n \geq 1$, besitzt ebenfalls die beiden Häufungspunkte $+1$ und -1 , denn es gilt

$$\lim_{k \rightarrow \infty} a_{2k} = \lim_{k \rightarrow \infty} \left(1 + \frac{1}{2k}\right) = 1$$

und analog $\lim_{k \rightarrow \infty} a_{2k+1} = -1$.

(5.3) Die Folge $a_n := n$, $n \in \mathbb{N}$, besitzt keinen Häufungspunkt, da jede Teilfolge unbeschränkt ist, also nicht konvergiert.

(5.4) Die Folge

$$a_n := \begin{cases} n & \text{falls } n \text{ gerade,} \\ \frac{1}{n} & \text{falls } n \text{ ungerade,} \end{cases}$$

ist unbeschränkt, besitzt aber den Häufungspunkt 0 , da die Teilfolge $(a_{2k+1})_{k \in \mathbb{N}}$ gegen 0 konvergiert.

(5.5) Für jede konvergente Folge ist der Limes ihr einziger Häufungspunkt.

Monotone Folgen

Definition. Eine Folge $(a_n)_{n \in \mathbb{N}}$ reeller Zahlen heißt

- i) *monoton wachsend*, falls $a_n \leq a_{n+1}$ für alle $n \in \mathbb{N}$,
- ii) *streng monoton wachsend*, falls $a_n < a_{n+1}$ für alle $n \in \mathbb{N}$,
- iii) *monoton fallend*, falls $a_n \geq a_{n+1}$ für alle $n \in \mathbb{N}$,
- iv) *streng monoton fallend*, falls $a_n > a_{n+1}$ für alle $n \in \mathbb{N}$.

Satz 7. Jede beschränkte monotone Folge (a_n) reeller Zahlen konvergiert.

Dies ist ein Konvergenzkriterium, das häufig angewendet werden kann, da in der Praxis viele Folgen monoton sind. Beispielsweise definiert jeder positive (negative) unendliche Dezimalbruch eine beschränkte, monoton wachsende (bzw. fallende) Folge.

Beweis. Nach dem Satz von Bolzano-Weierstraß besitzt die Folge (a_n) eine konvergente Teilfolge (a_{n_k}) . Sei a der Limes dieser Teilfolge. Wir zeigen, dass auch die gesamte Folge gegen a konvergiert. Dabei setzen wir voraus, dass die Folge (a_n) monoton wächst; für monoton fallende Folgen geht der Beweis analog.

Sei $\varepsilon > 0$ vorgegeben. Dann existiert ein $k_0 \in \mathbb{N}$, so dass

$$|a_{n_k} - a| < \varepsilon \quad \text{für alle } k \geq k_0.$$

Sei $N := n_{k_0}$. Zu jedem $n \geq N$ gibt es ein $k \geq k_0$ mit $n_k \leq n < n_{k+1}$. Da die Folge (a_n) monoton wächst, folgt daraus

$$a_{n_k} \leq a_n \leq a_{n_{k+1}} \leq a,$$

also

$$|a_n - a| \leq |a_{n_k} - a| < \varepsilon, \quad \text{q.e.d.}$$

Schluss-Bemerkung zu den Axiomen der reellen Zahlen. Mit den Körper-Axiomen, den Anordnungs-Axiomen, dem Archimedischen Axiom und dem Vollständigkeits-Axiom haben wir nun alle Axiome der reellen Zahlen aufgezählt. Ein Körper, in dem diese Axiome erfüllt sind, heißt vollständiger, archimedisch angeordneter Körper. Man kann beweisen, dass jeder vollständige, archimedisch angeordnete Körper dem Körper der reellen Zahlen isomorph

ist, dass also die genannten Axiome die reellen Zahlen vollständig charakterisieren.

Wir haben hier die reellen Zahlen als gegeben betrachtet. Man kann aber auch, ausgehend von den natürlichen Zahlen (die nach einem Ausspruch von L. Kronecker vom lieben Gott geschaffen worden sind, während alles andere Menschenwerk sei), nacheinander die ganzen Zahlen, die rationalen Zahlen und die reellen Zahlen konstruieren und dann die Axiome beweisen. Diesen Aufbau des Zahlensystems sollte jeder Mathematik-Student im Laufe seines Studiums kennenlernen. Wir verweisen hierzu auf die Literatur, z.B. [L], [Z].

ANHANG

Zur Darstellung reeller Zahlen im Computer

Zahlen werden in heutigen Computern meist *binär*, d.h. bzgl. der Basis 2 dargestellt. Natürlich ist es unmöglich, reelle Zahlen als unendliche 2-adische Brüche zu speichern, sondern man muss sich auf eine endliche Anzahl von Ziffern (*bits* = *binary digits*) beschränken. Um betragsmäßig große und kleine Zahlen mit derselben relativen Genauigkeit darzustellen, verwendet man eine sog. *Gleitpunkt*-Darstellung¹ der Form

$$x = \pm a_0 . a_1 a_2 a_3 \dots a_m \cdot 2^r,$$

wobei r ein ganzzahliger Exponent ist. Das Vorzeichen wird als $(-1)^s$ durch ein Bit $s \in \{0, 1\}$ dargestellt.

$$\xi := a_0 . a_1 a_2 a_3 \dots a_m := \sum_{\mu=0}^m a_\mu \cdot 2^{-\mu}, \quad a_\mu \in \{0, 1\},$$

ist die sog. *Mantisse*, die man für $x \neq 0$ durch geeignete Wahl des Exponenten im Bereich $1 \leq \xi < 2$ annehmen kann, was gleichbedeutend mit $a_0 = 1$ ist. Der Exponent r wird natürlich auch binär mit einer begrenzten Anzahl von Bits gespeichert. Um nicht das Vorzeichen von r eigens abspeichern zu müssen, schreibt man r in der Form $r = e - e_*$ mit einem festen Offset $e_* > 0$ und

$$e = \sum_{v=0}^{k-1} e_v \cdot 2^v \geq 0, \quad e_v \in \{0, 1\}.$$

Häufig werden insgesamt 64 Bits zur Darstellung einer reellen Zahl verwendet (Datentyp `DOUBLE PRECISION` in FORTRAN oder `double float` in

¹Statt Gleitpunkt sagt man auch Fließpunkt oder Fließkomma, engl. *floating point*.

den Programmiersprachen C, Java, usw.). Dabei wird üblicherweise der IEEE-Standard² befolgt, der hierfür $m = 52$, $k = 11$ und $e_* = 1023$ vorsieht³. Das Bit a_0 wird nicht gespeichert, sondern ist implizit gegeben. Insgesamt wird daher ein `double float` durch folgenden Bit-Vektor dargestellt:

$$(s, e_{10}, e_9, \dots, e_0, a_1, a_2, \dots, a_{52}) \in \{0, 1\}^{64}.$$

Der Exponent $e = \sum_{v=0}^{10} e_v \cdot 2^v$ kann Werte im Bereich $0 \leq e \leq 2^{11} - 1 = 2047$ annehmen. Falls $1 \leq e \leq 2046$, wird das implizite Bit $a_0 = 1$ gesetzt, es wird also die Zahl

$$x = (-1)^s 2^{e-1023} \left(1 + \sum_{\mu=1}^{52} a_\mu \cdot 2^{-\mu} \right)$$

dargestellt; für $e = 0$ wird vereinbart

$$x = (-1)^s 2^{-1022} \sum_{\mu=1}^{52} a_\mu 2^{-\mu},$$

während der Fall $e = 2047$ der Anzeige von Fehlerbedingungen vorbehalten ist. Die Zahl 0 wird also durch den Bit-Vektor, der aus lauter Nullen besteht, dargestellt. Die kleinste darstellbare positive Zahl ist danach $2^{-1074} \approx 4.94 \cdot 10^{-324}$, die größte Zahl $2^{1023}(2 - 2^{-52}) \approx 1.79 \cdot 10^{308}$. Die arithmetischen Operationen (Addition, Multiplikation, ...) auf Gleitpunktzahlen sind im Allgemeinen mit Fehlern versehen, da das exakte Resultat (falls es nicht überhaupt dem Betrag nach größer als die größte darstellbare Zahl ist, also zu Überlauf führt), noch auf eine mit der gegebenen Mantissenlänge verträgliche Zahl gerundet werden muss. Die Gleitpunkt-Arithmetik wird meist durch sog. mathematische Coprozessoren unterstützt, die z.B. im Falle der auf PCs weit verbreiteten Intel-Prozessoren intern mit 80-Bit-Zahlen arbeiten, wobei 64 Bits für die Mantisse, 15 Bits für den Exponenten und ein Vorzeichen-Bit verwendet werden. Beliebig einstellbare Genauigkeit wird meist nicht direkt durch die Hardware, sondern durch Software realisiert. Man vergesse aber nicht, dass die Gleitpunkt-Arithmetik inhärent fehlerbehaftet ist. Selbst so eine einfache Zahl wie $\frac{1}{10}$ wird binär auch bei noch so großer Mantissen-Länge nicht exakt dargestellt.

²IEEE = Institute of Electrical and Electronics Engineers

³Bei 32-bit floats sind die entsprechenden Zahlen $m = 23$, $k = 8$ und $e_* = 127$.

AUFGABEN

5.1. Man entwickle die Zahl $x = \frac{1}{7}$ in einen b -adischen Bruch für $b = 2, 7, 10, 16$. Im 16-adischen System (= Hexadezimalsystem) verwende man für die Ziffern 10 bis 15 die Buchstaben A bis F.

5.2. Ein b -adischer Bruch

$$a_{-k} \dots a_0 . a_1 a_2 a_3 a_4 \dots$$

heißt *periodisch*, wenn natürliche Zahlen $r, s \geq 1$ existieren, so dass

$$a_{n+s} = a_n \quad \text{für alle } n \geq r.$$

Man beweise: Ein b -adischer Bruch ist genau dann periodisch, wenn er eine rationale Zahl darstellt.

5.3. Gegeben seien zwei (unendliche) g -adische Brüche ($g \geq 2$),

$$0 . a_1 a_2 a_3 a_4 \dots,$$

$$0 . b_1 b_2 b_3 b_4 \dots,$$

die gegen dieselbe Zahl $x \in \mathbb{R}$ konvergieren. Man zeige: Entweder gilt $a_n = b_n$ für alle $n \geq 1$ oder es existiert eine natürliche Zahl $k \geq 1$, so dass (nach evtl. Vertauschung der Rollen von a und b) gilt:

$$\left\{ \begin{array}{ll} a_i = b_i & \text{für alle } i < k, \\ a_k = b_k + 1, \\ a_n = 0 & \text{für alle } n > k, \\ b_n = g - 1 & \text{für alle } n > k. \end{array} \right.$$

5.4. Man bestimme die 64-Bit-IEEE-Darstellung der Zahlen $z_n := 10^n$ für $n = 2, 1, 0, -1, -2$.

5.5. Es sei $Q_{64} \subset \mathbb{R}$ die Menge aller durch den 64-Bit-IEEE-Standard exakt dargestellten reellen Zahlen (diese sind natürlich alle rational) und R_{64} das Intervall $R_{64} := \{x \in \mathbb{R} : |x| < 2^{1024}\}$. Eine Abbildung

$$\rho : R_{64} \longrightarrow Q_{64}$$

werde wie folgt definiert: Für $x \in R_{64}$ sei $\rho(x)$ die Zahl aus Q_{64} , die von x den kleinsten Abstand hat. Falls zwei Elemente aus Q_{64} von x denselben Abstand haben, werde dasjenige gewählt, in deren IEEE-Darstellung das Bit $a_{52} = 0$ ist. Man überlege sich, dass dadurch ρ eindeutig definiert ist. Nunmehr werde

eine Addition

$$Q_{64} \times Q_{64} \longrightarrow Q_{64} \cup \{\diamond\}, \quad (x, y) \mapsto x \boxplus y,$$

durch folgende Vorschrift definiert: Falls $x + y \in R_{64}$, sei

$$x \boxplus y := \rho(x + y).$$

Falls aber $x + y \notin R_{64}$, setze man $x \boxplus y := \diamond$. Dabei sei $\diamond$ ein nicht zu $\mathbb{R}$ gehöriges Symbol, das als „undefiniert“ gelesen werde. (Seine Verwendung ist nur ein formaler Trick, damit $\boxplus$ ausnahmslos auf $Q_{64} \times Q_{64}$ definiert ist.)

a) Man zeige: Für alle $x, y \in Q_{64}$ gilt

$$(i) \ x \boxplus y = y \boxplus x, \quad (ii) \ x \boxplus 0 = x, \quad (iii) \ x \boxplus -x = 0.$$

b) Man zeige durch Angabe von Gegenbeispielen, dass das Assoziativ-Gesetz

$$(x \boxplus y) \boxplus z = x \boxplus (y \boxplus z)$$

in Q_{64} im Allgemeinen falsch ist, selbst wenn beide Seiten definiert sind.

Man gebe auch ein Beispiel von Zahlen $x, y, z \in Q_{64}$ an, so dass $x \boxplus y$, $(x \boxplus y) \boxplus z$ und $y \boxplus z$ alle zu Q_{64} gehören, aber $x \boxplus (y \boxplus z) = \diamond$.

5.6. Man zeige, dass $+1$ und -1 die einzigen Häufungspunkte der in den Beispielen (5.1) und (5.2) angegebenen Folgen sind.

5.7. Sei x eine vorgegebene reelle Zahl. Die Folge $(a_n(x))_{n \in \mathbb{N}}$ sei definiert durch

$$a_n(x) := nx - \lfloor nx \rfloor \quad \text{für alle } n \in \mathbb{N}.$$

Man beweise: Ist x rational, so hat die Folge nur endlich viele Häufungspunkte; ist x irrational, so ist jede reelle Zahl a mit $0 \leq a \leq 1$ Häufungspunkt der Folge $(a_n(x))_{n \in \mathbb{N}}$.

5.8. Man beweise: Eine Folge reeller Zahlen konvergiert dann und nur dann, wenn sie beschränkt ist und genau einen Häufungspunkt besitzt.

5.9. Man beweise: Jede Folge reeller Zahlen enthält eine monotone (wachsende oder fallende) Teilfolge.

5.10. Man zeige: Jede monoton wachsende (bzw. fallende) Folge $(a_n)_{n \in \mathbb{N}}$, die nicht konvergiert, divergiert bestimmt gegen $+\infty$ (bzw. $-\infty$).

5.11. Man beweise: Aus Satz 7 (jede beschränkte monotone Folge reeller Zahlen konvergiert) lässt sich das Archimedische Axiom und das Intervallschachtelungs-Prinzip ableiten (ohne das Vollständigkeits-Axiom zu benutzen).

§ 6 Wurzeln

In diesem Paragraphen beweisen wir als Anwendung des Vollständigkeits-Axioms die Existenz der Wurzeln positiver reeller Zahlen und geben gleichzeitig ein Iterationsverfahren zu ihrer Berechnung an. Dieses Verfahren, mit dem schon die Babylonier ihre Näherungswerte für die Quadratwurzeln der natürlichen Zahlen bestimmt haben sollen, konvergiert außerordentlich rasch und zählt auch noch heute im Computer-Zeitalter zu den effizientesten Algorithmen.

Sei $a > 0$ eine reelle Zahl, deren Quadratwurzel bestimmt werden soll. Wenn $x > 0$ Quadratwurzel von a ist, d.h. der Gleichung $x^2 = a$ genügt, gilt $x = \frac{a}{x}$, andernfalls ist $x \neq \frac{a}{x}$. Dann wird das arithmetische Mittel

$$x' := \frac{1}{2} \left(x + \frac{a}{x} \right)$$

ein besserer Näherungswert für die Wurzel sein und man kann hoffen, durch Wiederholung der Prozedur eine Folge zu erhalten, die gegen die Wurzel aus a konvergiert. Dass dies tatsächlich der Fall ist, beweisen wir jetzt.

Satz 1. *Seien $a > 0$ und $x_0 > 0$ reelle Zahlen. Die Folge $(x_n)_{n \in \mathbb{N}}$ sei durch*

$$x_{n+1} := \frac{1}{2} \left(x_n + \frac{a}{x_n} \right)$$

rekursiv definiert. Dann konvergiert die Folge (x_n) gegen die Quadratwurzel von a , d.h. gegen die eindeutig bestimmte positive Lösung der Gleichung $x^2 = a$.

Beweis. Wir gehen in mehreren Schritten vor.

1) Ein einfacher Beweis durch vollständige Induktion zeigt, dass $x_n > 0$ für alle $n \geq 0$, insbesondere die Division $\frac{a}{x_n}$ immer zulässig ist.

2) Es gilt $x_n^2 \geq a$ für alle $n \geq 1$, denn

$$\begin{aligned} x_n^2 - a &= \frac{1}{4} \left(x_{n-1} + \frac{a}{x_{n-1}} \right)^2 - a = \frac{1}{4} \left(x_{n-1}^2 + 2a + \frac{a^2}{x_{n-1}^2} \right) - a \\ &= \frac{1}{4} \left(x_{n-1} - \frac{a}{x_{n-1}} \right)^2 \geq 0. \end{aligned}$$

3) Es gilt $x_{n+1} \leq x_n$ für alle $n \geq 1$, denn

$$x_n - x_{n+1} = x_n - \frac{1}{2} \left(x_n + \frac{a}{x_n} \right) = \frac{1}{2x_n} (x_n^2 - a) \geq 0.$$

4) Wegen 3) ist $(x_n)_{n \geq 1}$ eine monoton fallende Folge, die durch 0 nach unten beschränkt ist, also nach §5, Satz 7 konvergiert. (Hier geht das Vollständigkeits-Axiom ein, denn es wurde beim Beweis des Satzes über die Konvergenz beschränkter monotoner Folgen benötigt.) Für den Grenzwert $x := \lim_{n \rightarrow \infty} x_n$ gilt nach §4, Corollar zu Satz 5, dass $x \geq 0$. Um den Wert von x zu bestimmen, benutzen wir die Gleichung

$$2x_{n+1}x_n = x_n^2 + a,$$

die sich aus der Rekursionsformel durch Multiplikation mit $2x_n$ ergibt. Da $\lim_{n \rightarrow \infty} x_{n+1} = \lim_{n \rightarrow \infty} x_n = x$, folgt aus den Regeln über das Rechnen mit Grenzwerten (§4, Satz 3)

$$2x^2 = x^2 + a,$$

also $x^2 = a$. Damit ist gezeigt, dass die Folge (x_n) gegen eine Quadratwurzel von a konvergiert.

5) Es ist noch die Eindeutigkeit zu zeigen. Sei y eine weitere positive Lösung der Gleichung $y^2 = a$. Dann ist

$$0 = x^2 - y^2 = (x + y)(x - y).$$

Da $x + y > 0$, muss $x - y = 0$ sein, also $x = y$, q.e.d.

Bezeichnung. Für eine reelle Zahl $a \geq 0$ wird die eindeutig bestimmte nicht-negative Lösung der Gleichung $x^2 = a$ mit

$$\sqrt{a} \quad \text{oder} \quad \text{sqrt}(a)$$

bezeichnet.

Bemerkung. Die Gleichung $x^2 = a$ hat für $a = 0$ nur die Lösung $x = 0$ und für $a > 0$ genau zwei Lösungen, nämlich $\sqrt{a}$ und $-\sqrt{a}$. Denn für jedes $x \in \mathbb{R}$ mit $x^2 = a$ gilt

$$(x + \sqrt{a})(x - \sqrt{a}) = x^2 - a = 0,$$

also muss einer der beiden Faktoren gleich 0 sein, d.h. $x = \pm\sqrt{a}$. Für $a < 0$ hat die Gleichung natürlich keine reelle Lösung, weil für jedes $x \in \mathbb{R}$ gilt $x^2 \geq 0$.

Numerisches Beispiel

Zur numerischen Rechnung ist es günstig, noch die Größe $y_n := a/x_n$ einzuführen. Dann ist $x_{n+1} = \frac{1}{2}(x_n + y_n)$. Für $n \geq 1$ gilt nach 2) die Abschätzung $x_n^2 \geq a$, also $x_n \geq \sqrt{a}$. Daraus folgt $y_n \leq \sqrt{a}$; man hat also

$$y_n \leq \sqrt{a} \leq x_n \quad \text{für alle } n \geq 1.$$

Zur Illustration des Algorithmus rechnen wir ein Beispiel mit dem Multipräsizisions-Interpreter ARIBAS. Durch den Befehl

```
==> set_floatprec(long_float) .
- : 128
```

wird die Rechengenauigkeit auf `long_float` eingestellt, d.h. reelle Zahlen werden von ARIBAS mit einer 128-bit Mantisse dargestellt (relative Genauigkeit 2^{-128}). Wir wollen die Quadratwurzel aus $a := 2$ berechnen und wählen $x_0 = 2$ und $y_0 = a/x_0 = 1$. Es werden die Werte x_n und $y_n = a/x_n$ für $n = 1, \dots, 6$ berechnet.

```
==> a := 2;
      x := a; y := 1;
      for n := 1 to 6 do
        x := (x + y)/2; y := a/x;
        writeln(n, " ");
        writeln(y); writeln(x);
      end.
```

Die Variablen `x` und `y` enthalten vor dem Eintritt in den n -ten Durchlauf der `for`-Schleife die Werte x_{n-1} und y_{n-1} ; diese werden dann durch x_n und y_n ersetzt und mit `writeln` ausgegeben. Insgesamt erhält man folgende Ausgabe:

```
1)
1.33333_33333_33333_33333_33333_33333_33333_33333_3
1.50000_00000_00000_00000_00000_00000_00000_00000_0
2)
1.41176_47058_82352_94117_64705_88235_29411_8
1.41666_66666_66666_66666_66666_66666_66666_66666_7
3)
1.41421_14384_74870_01733_10225_30329_28942_8
1.41421_56862_74509_80392_15686_27450_98039_2
4)
1.41421_35623_71500_18697_70836_68114_92557_7
1.41421_35623_74689_91062_62955_78890_13491_0
5)
1.41421_35623_73095_04880_16878_24916_86591_4
1.41421_35623_73095_04880_16896_23502_53024_4
6)
1.41421_35623_73095_04880_16887_24209_69807_9
1.41421_35623_73095_04880_16887_24209_69807_9
```

Der Wert von $\sqrt{2}$ liegt nach dem oben Gesagten stets zwischen den unmittelbar untereinander stehenden Zahlen. Man kann gut beobachten, wie die Anzahl der übereinstimmenden Dezimalstellen mit jedem Schritt steigt. Bereits nach 6 Iterations-Schritten ist die Wurzel aus 2 auf 36 Dezimalstellen genau bestimmt. (Allerdings kann sich die letzte berechnete Stelle bei Erhöhung der Genauigkeit noch ändern; tatsächlich ergibt sich statt der letzten 9 genauer 85696...)

Geschwindigkeit der Konvergenz

Die in dem Beispiel sichtbare schnelle Konvergenz wollen wir nun im allgemeinen Fall untersuchen. Dazu definieren wir den relativen Fehler f_n im n -ten Iterationsschritt durch die Gleichung

$$x_n = \sqrt{a}(1 + f_n).$$

Es ist $f_n \geq 0$ für $n \geq 1$. Einsetzen in die Gleichung $x_{n+1} = \frac{1}{2}(x_n + \frac{a}{x_n})$ ergibt nach Kürzung durch $\sqrt{a}$

$$1 + f_{n+1} = \frac{1}{2} \left(1 + f_n + \frac{1}{1 + f_n} \right).$$

Daraus folgt

$$f_{n+1} = \frac{1}{2} \cdot \frac{f_n^2}{1 + f_n} \leq \frac{1}{2} \min(f_n, f_n^2).$$

Da Zahlen im Computer meist binär, d.h. bzgl. der Basis 2 dargestellt werden, ist die Multiplikation mit ganzzahligen Potenzen von 2 trivial. So kann man jede positive reelle Zahl a leicht in die Form $a = 2^{2k}a_0$ mit $k \in \mathbb{Z}$, $1 \leq a_0 < 4$, bringen. Es ist dann $\sqrt{a} = 2^k \sqrt{a_0}$; also kann man ohne Beschränkung der Allgemeinheit $1 \leq a < 4$ voraussetzen. Wählt man dann $x_0 = a$, so ist $\sqrt{a} \leq x_0 < 2\sqrt{a}$, d.h. $0 \leq f_0 < 1$. Mit der obigen Rekursionsformel für den relativen Fehler ergibt sich $f_1 < 1/4$, $f_2 < 1/40$, ..., $f_5 < 1.2 \cdot 10^{-15}$, $f_6 < 10^{-30}$ etc. Die Zahl der gültigen Dezimalstellen verdoppelt sich also mit jedem Schritt. Man spricht von *quadratischer* Konvergenz.

Der angegebene Algorithmus zur Wurzelberechnung hat neben seiner schnellen Konvergenz noch den Vorteil, *selbstkorrigierend* zu sein. Denn da der Anfangswert $x_0 > 0$ beliebig vorgegeben werden kann, beginnt nach eventuellen Rechen-, insbesondere Rundungsfehlern, der Algorithmus eben wieder mit dem fehlerhaften Wert von x_n statt mit x_0 . Wollten wir etwa $\sqrt{2}$ auf 100 Dezimalstellen genau berechnen, so müssten wir nicht die Rechnung von Anfang an mit 100-stelliger Genauigkeit wiederholen, sondern könnten mit dem erhal-

tenen 36-stelligen Näherungswert beginnen und erhielten nach zwei weiteren Schritten das Ergebnis.

Es sei jedoch bemerkt, dass die Verhältnisse nicht immer so günstig liegen. Bei vielen Näherungs-Verfahren der numerischen Mathematik ist die Fehlerabschätzung viel schwieriger; Rundungsfehler können sich akkumulieren und aufschaukeln, wodurch manchmal sogar die Konvergenz, die unter der Prämisse der exakten Rechnung bewiesen worden ist, gefährdet wird.

***k*-te Wurzeln**

Die für die Quadratwurzeln angestellten Überlegungen lassen sich leicht auf *k*-te Wurzeln verallgemeinern.

Satz 2. *Sei $k \geq 2$ eine natürliche Zahl und $a > 0$ eine positive reelle Zahl. Dann konvergiert für jeden Anfangswert $x_0 > 0$ die durch*

$$x_{n+1} := \frac{1}{k} \left((k-1)x_n + \frac{a}{x_n^{k-1}} \right)$$

rekursiv definierte Folge $(x_n)_{n \in \mathbb{N}}$ gegen die eindeutig bestimmte positive Lösung der Gleichung $x^k = a$.

Bezeichnung. Diese Lösung heißt *k*-te Wurzel von *a* und wird mit $\sqrt[k]{a}$ bezeichnet. Im Spezialfall $k = 2$ schreibt man kürzer $\sqrt{a}$ statt $\sqrt[2]{a}$.

Beweis. a) Dass es nicht mehr als eine positive Lösung geben kann, ergibt sich daraus, dass aus $0 \leq z_1 < z_2$ folgt $z_1^k < z_2^k$.

b) Wir zeigen jetzt, dass der Grenzwert $x := \lim_{n \rightarrow \infty} x_n$ im Falle der Existenz die Gleichung $x^k = a$ erfüllt. Multiplikation der Rekursionsgleichung mit kx_n^{k-1} ergibt nämlich

$$kx_{n+1}x_n^{k-1} = (k-1)x_n^k + a,$$

woraus durch Grenzübergang $n \rightarrow \infty$ folgt

$$kx^k = (k-1)x^k + a,$$

d.h. $x^k = a$.

c) Es bleibt nur noch die Konvergenz zu beweisen. Dazu formen wir die Rekursionsformel wie folgt um:

$$x_{n+1} = \frac{1}{k}x_n \left((k-1) + \frac{a}{x_n^k} \right) = x_n \left(1 + \frac{1}{k} \left(\frac{a}{x_n^k} - 1 \right) \right) \quad (*)$$

Aus der Bernoullischen Ungleichung $(1 + \xi)^k \geq 1 + k\xi$ für $\xi \geq -1$ folgt

$$\left(1 + \frac{1}{k} \left(\frac{a}{x_n^k} - 1\right)\right)^k \geq 1 + \left(\frac{a}{x_n^k} - 1\right) = \frac{a}{x_n^k},$$

also $x_{n+1}^k \geq a$. Deshalb ist $a/x_n^k \leq 1$ für alle $n \geq 1$, d.h.

$$\frac{1}{k} \left(\frac{a}{x_n^k} - 1\right) \leq 0$$

und man liest aus (*) ab $x_{n+1} \leq x_n$ für $n \geq 1$. Die Folge $(x_n)_{n \geq 1}$ ist also monoton fallend und durch 0 nach unten beschränkt, konvergiert also gegen eine reelle Zahl $x \geq 0$, für die nach b) gilt $x^k = a$, q.e.d.

Bemerkungen. 1) Natürlich enthält Satz 2 als Spezialfall den Satz 1. Man kann sich fragen, warum beim Iterationsschritt für die k -te Wurzel nicht die Formel $x_{n+1} = \frac{1}{2}(x_n + a/x_n^{k-1})$ verwendet wird. Den tieferen Grund dafür werden wir in §17 kennenlernen, wo das Iterationsverfahren zum Wurzelziehen sich als Spezialfall eines viel allgemeineren Approximations-Verfahrens von Newton herausstellen wird.

2) Ein weiterer Beweis für die Existenz der Wurzeln wird in §12 gegeben, als Anwendung eines allgemeinen Satzes über Umkehrfunktionen.

AUFGABEN

6.1. Man beweise für $a \geq 0$, $b \geq 0$ die Ungleichung zwischen geometrischem und arithmetischem Mittel

$$\sqrt{ab} \leq \frac{1}{2}(a+b),$$

wobei Gleichheit genau dann eintritt, wenn $a = b$.

6.2. Seien $a \geq 0$, $b \geq 0$ reelle Zahlen. Die Folgen $(a_n)_{n \in \mathbb{N}}$, $(b_n)_{n \in \mathbb{N}}$ seien rekursiv definiert durch $a_0 := a$, $b_0 := b$ und

$$a_{n+1} := \sqrt{a_n b_n}, \quad b_{n+1} := \frac{1}{2}(a_n + b_n) \quad \text{für alle } n \in \mathbb{N}.$$

i) Man zeige, dass beide Folgen konvergieren und denselben Grenzwert haben. Dieser Grenzwert werde mit $M(a, b)$ bezeichnet und heißt das *arithmetisch-geometrische Mittel* von a und b .

ii) Man beweise für $a > 0$, $b \geq 0$ die Beziehung

$$M(a, b) = aM(1, b/a).$$

iii) Man berechne $M(1, 2)$ mit einem Fehler $< 10^{-6}$.

6.3. Beim Iterations-Verfahren

$$x_0 > 0, \quad x_{n+1} := \frac{1}{3} \left(2x_n + \frac{a}{x_n^2} \right)$$

zur Berechnung der 3. Wurzel einer positiven Zahl $a > 0$ definiere man den n -ten relativen Fehler f_n durch

$$x_n = \sqrt[3]{a}(1 + f_n).$$

Man leite eine Rekursionsformel für die Folge (f_n) her und beweise

$$f_{n+1} \leq f_n^2 \quad \text{für alle } n \geq 1.$$

6.4. Seien $a > 0$ und $x_0 > 0$ reelle Zahlen mit $ax_0 < 2$. Die Folge $(x_n)_{n \in \mathbb{N}}$ werde rekursiv definiert durch

$$x_{n+1} := x_n(1 + \varepsilon_n), \quad \text{wobei } \varepsilon_n := 1 - ax_n.$$

Man beweise, dass die Folge (x_n) gegen $1/a$ konvergiert.

Anleitung. Man zeige dazu: $\varepsilon_{n+1} = \varepsilon_n^2$ für alle $n \geq 0$.

Bemerkung. Dieser Algorithmus kann benutzt werden, um die Division auf die Multiplikation zurückzuführen.

6.5. Man berechne

$$\sqrt{1 + \sqrt{1 + \sqrt{1 + \sqrt{1 + \dots}}}},$$

d.h. den Limes der Folge $(a_n)_{n \in \mathbb{N}}$ mit $a_0 := 1$ und $a_{n+1} := \sqrt{1 + a_n}$.

6.6. Der Wert des unendlichen Kettenbruchs

$$1 + \frac{1}{1 + \frac{1}{1 + \frac{1}{1 + \frac{1}{1 + \dots}}}}$$

ist definiert als der Limes der Folge $(a_n)_{n \in \mathbb{N}}$ mit $a_0 := 1$ und $a_{n+1} := 1 + \frac{1}{a_n}$.

a) Man zeige $a_{n-1} = \frac{f_{n+1}}{f_n}$ für alle $n \geq 1$, wobei f_n die in (4.6) definierten Fibonacci-Zahlen sind.

b) Man beweise $\lim_{n \rightarrow \infty} a_n = \frac{1 + \sqrt{5}}{2}$.

Bemerkung. Der Limes ist der berühmte *goldene Schnitt*, der durch

$$g : 1 = 1 : (g - 1), \quad g > 1,$$

definiert ist.

6.7. Die drei Folgen $(a_n)_{n \in \mathbb{N}}$, $(b_n)_{n \in \mathbb{N}}$, $(c_n)_{n \in \mathbb{N}}$ seien definiert durch

$$a_n := \sqrt{n+1000} - \sqrt{n},$$

$$b_n := \sqrt{n + \sqrt{n}} - \sqrt{n},$$

$$c_n := \sqrt{n + \frac{n}{1000}} - \sqrt{n}.$$

Man zeige: Für alle $n < 10^6$ gilt $a_n > b_n > c_n$, aber

$$\lim_{n \rightarrow \infty} a_n = 0, \quad \lim_{n \rightarrow \infty} b_n = \frac{1}{2}, \quad \lim_{n \rightarrow \infty} c_n = \infty.$$

6.8. Man beweise:

a) $\lim_{n \rightarrow \infty} \left(\sqrt[3]{n + \sqrt{n}} - \sqrt[3]{n} \right) = 0.$

b) $\lim_{n \rightarrow \infty} \left(\sqrt[3]{n + \sqrt[3]{n^2}} - \sqrt[3]{n} \right) = \frac{1}{3}.$

6.9. a) Man zeige: Für alle natürlichen Zahlen $n \geq 1$ gilt

$$\sqrt[n]{n} \leq 1 + \frac{2}{\sqrt{n}}.$$

Anleitung. Man verwende dazu Aufgabe 3.4.

b) Man folgere aus Teil a)

$$\lim_{n \rightarrow \infty} \sqrt[n]{n} = 1.$$

6.10. Für eine reelle Zahl $x > 0$ und eine rationale Zahl $r = m/n$, ($m, n \in \mathbb{Z}$, $n > 0$) definiert man

$$x^r := \sqrt[n]{x^m}.$$

a) Man zeige, dass diese Definition unabhängig von der Darstellung von r ist, d.h. aus $m/n = \ell/k$ folgt $\sqrt[n]{x^m} = \sqrt[k]{x^\ell}$.

b) Für $x, y \in \mathbb{R}_+^*$ und $r, s \in \mathbb{Q}$ beweise man die Rechenregeln

$$x^r y^r = (xy)^r, \quad x^r x^s = x^{r+s}, \quad (x^r)^s = x^{rs}.$$

§ 7 Konvergenz-Kriterien für Reihen

In diesem Paragraphen beweisen wir die wichtigsten Konvergenz-Kriterien für unendliche Reihen und behandeln einige typische Beispiele.

Wendet man das Vollständigkeits-Axiom über die Konvergenz von Cauchy-Folgen auf Reihen an, so erhält man folgendes Kriterium.

Satz 1 (Cauchysches Konvergenz-Kriterium). *Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge reeller Zahlen. Die Reihe $\sum_{n=0}^{\infty} a_n$ konvergiert genau dann, wenn gilt:*

Zu jedem $\varepsilon > 0$ existiert ein $N \in \mathbb{N}$, so dass

$$\left| \sum_{k=m}^n a_k \right| < \varepsilon \quad \text{für alle } n \geq m \geq N.$$

Beweis. Wir bezeichnen mit $S_N := \sum_{k=0}^N a_k$ die N -te Partialsumme. Dann ist

$$S_n - S_{m-1} = \sum_{k=m}^n a_k.$$

Die angegebene Bedingung drückt deshalb einfach aus, dass die Folge (S_n) der Partialsummen eine Cauchy-Folge ist, was gleichbedeutend mit ihrer Konvergenz ist.

Bemerkung. Aus Satz 1 folgt unmittelbar: Das Konvergenzverhalten einer Reihe ändert sich nicht, wenn man endlich viele Summanden abändert. (Nur die Summe ändert sich.)

Satz 2. *Eine notwendige (aber nicht hinreichende) Bedingung für die Konvergenz einer Reihe $\sum_{n=0}^{\infty} a_n$ ist, dass*

$$\lim_{n \rightarrow \infty} a_n = 0.$$

Beweis. Wenn die Reihe konvergiert, gibt es nach Satz 1 zu vorgegebenem $\varepsilon > 0$ ein $N \in \mathbb{N}$, so dass

$$\left| \sum_{k=m}^n a_k \right| < \varepsilon \quad \text{für alle } n \geq m \geq N.$$

Insbesondere gilt daher (für $n = m$)

$$|a_n| < \varepsilon \quad \text{für alle } n \geq N.$$

Daraus folgt $\lim a_n = 0$, q.e.d.

Beispielsweise divergiert die Reihe $\sum_{n=0}^{\infty} (-1)^n$, da die Reihenglieder nicht gegen 0 konvergieren. Ein Beispiel dafür, dass die Bedingung $\lim a_n = 0$ für die Konvergenz nicht ausreicht, behandeln wir in (7.1) im Anschluss an den nächsten Satz.

Satz 3. Eine Reihe $\sum_{n=0}^{\infty} a_n$ mit $a_n \geq 0$ für alle $n \in \mathbb{N}$ konvergiert genau dann, wenn die Reihe (d.h. die Folge der Partialsummen) beschränkt ist.

Beweis. Da $a_n \geq 0$, ist die Folge der Partialsummen

$$S_n = \sum_{k=0}^n a_k, \quad n \in \mathbb{N},$$

monoton wachsend. Die Behauptung folgt deshalb aus dem Satz über die Konvergenz monotoner beschränkter Folgen.

Beispiele

(7.1) Die *harmonische Reihe* $\sum_{n=1}^{\infty} \frac{1}{n}$.

Die Reihenglieder konvergieren gegen 0, trotzdem divergiert die Reihe. Dazu betrachten wir die speziellen Partialsummen

$$\begin{aligned} S_{2^k} &= \sum_{n=1}^{2^k} \frac{1}{n} = 1 + \frac{1}{2} + \sum_{i=1}^{k-1} \left(\sum_{n=2^i+1}^{2^{i+1}} \frac{1}{n} \right) \\ &= 1 + \frac{1}{2} + \left(\frac{1}{3} + \frac{1}{4} \right) \\ &\quad + \left(\frac{1}{5} + \frac{1}{6} + \frac{1}{7} + \frac{1}{8} \right) + \dots + \left(\frac{1}{2^{k-1}+1} + \dots + \frac{1}{2^k} \right). \end{aligned}$$

Da die Summe jeder Klammer $\geq \frac{1}{2}$ ist, folgt

$$S_{2^k} \geq 1 + \frac{k}{2}.$$

Also ist die Folge der Partialsummen unbeschränkt, d.h. es gilt

$$\sum_{n=1}^{\infty} \frac{1}{n} = \infty.$$

(7.2) Die Reihen $\sum_{n=1}^{\infty} \frac{1}{n^k}$ für $k > 1$.

Wir beweisen, dass diese Reihen konvergieren, indem wir zeigen, dass die Partialsummen durch $\frac{1}{1-2^{-k+1}}$ beschränkt sind. Zu beliebigem $N \in \mathbb{N}$ gibt es ein $m \in \mathbb{N}$ mit $N \leq 2^{m+1} - 1$. Damit gilt

$$\begin{aligned} S_N &\leq \sum_{n=1}^{2^{m+1}-1} \frac{1}{n^k} = 1 + \left(\frac{1}{2^k} + \frac{1}{3^k}\right) + \dots + \left(\sum_{n=2^m}^{2^{m+1}-1} \frac{1}{n^k}\right) \\ &\leq \sum_{i=0}^m 2^i \frac{1}{(2^i)^k} = \sum_{i=0}^m \left(\frac{1}{2^{k-1}}\right)^i \\ &\leq \sum_{i=0}^{\infty} (2^{-k+1})^i = \frac{1}{1-2^{-k+1}}, \quad \text{q.e.d.} \end{aligned}$$

Bemerkungen. a) Die hier zur Untersuchung der Konvergenz der Reihen $\sum \frac{1}{n^k}$ benutzte Methode ist ein Spezialfall des sog. Reihenverdichtungs-Kriteriums, siehe Aufgabe 7.3.

b) Für alle geraden ganzen Zahlen $k \geq 2$ gibt es explizite Formeln für die Limiten der Reihen $\sum_{n=1}^{\infty} \frac{1}{n^k}$. Z.B. gilt

$$\sum_{n=1}^{\infty} \frac{1}{n^2} = \frac{\pi^2}{6}, \quad \sum_{n=1}^{\infty} \frac{1}{n^4} = \frac{\pi^4}{90}, \quad \sum_{n=1}^{\infty} \frac{1}{n^6} = \frac{\pi^6}{945},$$

siehe dazu (21.9) und §22, Satz 11.

Während sich Satz 3 auf Reihen mit lauter nicht-negativen Gliedern bezog, behandeln wir jetzt ein Konvergenz-Kriterium für *alternierende* Reihen, das sind Reihen, deren Glieder abwechselndes Vorzeichen haben.

Satz 4 (Leibniz'sches Konvergenz-Kriterium). *Sei $(a_n)_{n \in \mathbb{N}}$ eine monoton fallende Folge nicht-negativer Zahlen mit $\lim_{n \rightarrow \infty} a_n = 0$. Dann konvergiert die alternierende Reihe*

$$\sum_{n=0}^{\infty} (-1)^n a_n.$$

Beweis. Wir setzen $s_k := \sum_{n=0}^k (-1)^n a_n$. Da $s_{2k+2} - s_{2k} = -a_{2k+1} + a_{2k+2} \leq 0$, gilt

$$s_0 \geq s_2 \geq s_4 \geq \dots \geq s_{2k} \geq s_{2k+2} \geq \dots$$

Entsprechend ist wegen $s_{2k+3} - s_{2k+1} = a_{2k+2} - a_{2k+3} \geq 0$

$$s_1 \leq s_3 \leq s_5 \leq \dots \leq s_{2k+1} \leq s_{2k+3} \leq \dots$$

Außerdem gilt wegen $s_{2k+1} - s_{2k} = -a_{2k+1} \leq 0$

$$s_{2k+1} \leq s_{2k} \quad \text{für alle } k \in \mathbb{N}.$$

Die Folge $(s_{2k})_{k \in \mathbb{N}}$ ist also monoton fallend und beschränkt, da $s_{2k} \geq s_1$ für alle k . Nach §5, Satz 7, existiert daher der Limes

$$\lim_{k \rightarrow \infty} s_{2k} =: S.$$

Analog ist $(s_{2k+1})_{k \in \mathbb{N}}$ monoton wachsend und beschränkt, also existiert

$$\lim_{k \rightarrow \infty} s_{2k+1} =: S'.$$

Wir zeigen nun, dass $S = S'$ und dass die gesamte Folge $(s_n)_{n \in \mathbb{N}}$ gegen S konvergiert. Zunächst ist

$$S - S' = \lim_{k \rightarrow \infty} (s_{2k} - s_{2k+1}) = \lim_{k \rightarrow \infty} a_{2k+1} = 0.$$

Sei nun $\varepsilon > 0$ vorgegeben. Dann gibt es $N_1, N_2 \in \mathbb{N}$, so dass

$$|s_{2k} - S| < \varepsilon \quad \text{für } k \geq N_1 \quad \text{und} \quad |s_{2k+1} - S| < \varepsilon \quad \text{für } k \geq N_2.$$

Wir setzen $N := \max(2N_1, 2N_2 + 1)$. Dann gilt

$$|s_n - S| < \varepsilon \quad \text{für alle } n \geq N, \quad \text{q.e.d.}$$

Das Konvergenzverhalten der alternierenden Reihen lässt sich durch Bild 7.1 veranschaulichen.

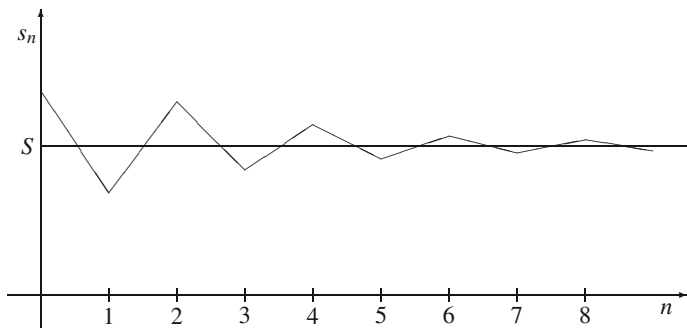


Bild 7.1 Zum Leibniz'schen Konvergenz-Kriterium

Beispiele

(7.3) Die *alternierende harmonische Reihe* $\sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{n}$ konvergiert nach dem Leibniz'schen Konvergenz-Kriterium. Wir werden in §22 sehen, dass

$$1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \frac{1}{5} - \frac{1}{6} \pm \dots = \log 2.$$

Dabei ist $\log 2 = 0.69314718\dots$ der natürliche Logarithmus von 2.

(7.4) Ebenso konvergiert die *Leibniz'sche Reihe* $\sum_{k=0}^{\infty} \frac{(-1)^k}{2k+1}$. Für sie zeigte Leibniz, dass

$$1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} + \frac{1}{9} - \frac{1}{11} \pm \dots = \frac{\pi}{4}.$$

Wir werden dies in §22 beweisen.

Absolute Konvergenz

Definition. Eine Reihe $\sum_{n=0}^{\infty} a_n$ heißt *absolut konvergent*, falls die Reihe der Absolutbeträge $\sum_{n=0}^{\infty} |a_n|$ konvergiert.

Bemerkung. Da die Partialsummen der Reihe $\sum |a_n|$ monoton wachsen, gilt nach §5, Satz 7: Eine Reihe $\sum a_n$ ist genau dann absolut konvergent, wenn

$$\sum_{n=0}^{\infty} |a_n| < \infty.$$

Satz 5. Eine absolut konvergente Reihe konvergiert auch im gewöhnlichen Sinn.

Bemerkung. Wie das Beispiel der alternierenden harmonischen Reihe (7.3) zeigt, gilt die Umkehrung von Satz 5 nicht. Die absolute Konvergenz ist also eine schärfere Bedingung als die gewöhnliche Konvergenz.

Beweis. Sei $\sum_{n=0}^{\infty} a_n$ eine absolut konvergente Reihe. Nach dem Cauchyschen Konvergenz-Kriterium (Satz 1) für die Reihe $\sum |a_n|$ gibt es zu jedem $\varepsilon > 0$ ein $N \in \mathbb{N}$, so dass

$$\sum_{k=m}^n |a_k| < \varepsilon \quad \text{für alle } n \geq m \geq N.$$

Daraus folgt

$$\left| \sum_{k=m}^n a_k \right| \leq \sum_{k=m}^n |a_k| < \varepsilon \quad \text{für alle } n \geq m \geq N.$$

Wiederum nach dem Cauchyschen Konvergenz-Kriterium konvergiert daher $\sum_{n=0}^{\infty} a_n$, q.e.d.

Satz 6 (Majoranten-Kriterium). Sei $\sum_{n=0}^{\infty} c_n$ eine konvergente Reihe mit lauter nicht-negativen Gliedern und $(a_n)_{n \in \mathbb{N}}$ eine Folge mit

$$|a_n| \leq c_n \quad \text{für alle } n \in \mathbb{N}.$$

Dann konvergiert die Reihe $\sum_{n=0}^{\infty} a_n$ absolut.

Bezeichnung. Man nennt dann $\sum c_n$ eine Majorante von $\sum a_n$.

Beweis. Zu vorgegebenem $\varepsilon > 0$ existiert ein $N \in \mathbb{N}$, so dass

$$\left| \sum_{k=m}^n c_k \right| < \varepsilon \quad \text{für alle } n \geq m \geq N.$$

Daher ist

$$\sum_{k=m}^n |a_k| \leq \sum_{k=m}^n c_k < \varepsilon \quad \text{für alle } n \geq m \geq N.$$

Die Reihe $\sum |a_n|$ erfüllt also das Cauchysche Konvergenz-Kriterium, q.e.d.

Beispiel

(7.5) Wir beweisen noch einmal die Konvergenz der Reihen $\sum_{n=1}^{\infty} \frac{1}{n^k}$, $k \geq 2$, mithilfe des Majoranten-Kriteriums.

Nach Beispiel (4.9) konvergiert die Reihe $\sum_{n=1}^{\infty} \frac{1}{n(n+1)}$, also auch die Reihe $\sum_{n=1}^{\infty} \frac{2}{n(n+1)}$. Für $k \geq 2$ und alle $n \geq 1$ gilt

$$\frac{1}{n^k} \leq \frac{1}{n^2} \leq \frac{2}{n(n+1)},$$

daher ist $\sum \frac{2}{n(n+1)}$ Majorante von $\sum \frac{1}{n^k}$, q.e.d.

Hinweis. Wir werden später (in §20) noch ein sehr nützliches, dem Majoranten-Kriterium verwandtes Konvergenz-Kriterium kennenlernen, das Integralvergleichs-Kriterium. Mit diesem lassen sich die Reihen $\sum \frac{1}{n^k}$ besonders elegant behandeln.

Bemerkung. Satz 6 impliziert folgendes Divergenz-Kriterium:

Sei $\sum_{n=0}^{\infty} c_n$ eine divergente Reihe mit lauter nicht-negativen Gliedern und $(a_n)_{n \in \mathbb{N}}$ eine Folge mit $a_n \geq c_n$ für alle n . Dann divergiert auch die Reihe $\sum_{n=0}^{\infty} a_n$.

Denn andernfalls wäre $\sum a_n$ eine konvergente Majorante von $\sum c_n$, also müsste auch $\sum c_n$ konvergieren.

Satz 7 (Quotienten-Kriterium). Sei $\sum_{n=0}^{\infty} a_n$ eine Reihe mit $a_n \neq 0$ für alle $n \geq n_0$. Es gebe eine reelle Zahl θ mit $0 < \theta < 1$, so dass

$$\left| \frac{a_{n+1}}{a_n} \right| \leq \theta \quad \text{für alle } n \geq n_0.$$

Dann konvergiert die Reihe $\sum a_n$ absolut.

Beweis. Da ein Abändern endlich vieler Summanden das Konvergenzverhalten nicht ändert, können wir ohne Beschränkung der Allgemeinheit annehmen, dass

$$\left| \frac{a_{n+1}}{a_n} \right| \leq \theta \quad \text{für alle } n \in \mathbb{N}.$$

Daraus ergibt sich mit vollständiger Induktion

$$|a_n| \leq |a_0| \theta^n \quad \text{für alle } n \in \mathbb{N}.$$

Die Reihe $\sum_{n=0}^{\infty} |a_0| \theta^n$ ist daher Majorante von $\sum a_n$. Da

$$\sum_{n=0}^{\infty} |a_0| \theta^n = |a_0| \sum_{n=0}^{\infty} \theta^n = \frac{|a_0|}{1 - \theta}$$

konvergiert (geometrische Reihe), folgt aus dem Majoranten-Kriterium die Behauptung.

Bemerkung. Ein dem Quotienten-Kriterium verwandtes Kobvergenz-Kriterium ist das Wurzel-Kriterium, siehe Aufgabe 7.5.

Beispiele

(7.6) Wir beweisen die Konvergenz der Reihe $\sum_{n=1}^{\infty} \frac{n^2}{2^n}$.

Mit $a_n := \frac{n^2}{2^n}$ gilt für alle $n \geq 3$

$$\left| \frac{a_{n+1}}{a_n} \right| = \frac{(n+1)^2 2^n}{2^{n+1} n^2} = \frac{1}{2} \left(1 + \frac{1}{n} \right)^2$$

$$\leq \frac{1}{2} \left(1 + \frac{1}{3}\right)^2 = \frac{8}{9} =: \theta < 1,$$

das Quotienten-Kriterium ist also erfüllt.

(7.7) Man beachte, dass die Bedingung im Quotienten-Kriterium *nicht* lautet

$$\left| \frac{a_{n+1}}{a_n} \right| < 1 \quad \text{für alle } n \geq n_0, \quad (*)$$

sondern

$$\left| \frac{a_{n+1}}{a_n} \right| \leq \theta \quad \text{für alle } n \geq n_0$$

mit einem von n *unabhängigen* $\theta < 1$. Die Quotienten $\left| \frac{a_{n+1}}{a_n} \right|$ dürfen also nicht beliebig nahe an 1 herankommen. Dass die Bedingung (*) nicht ausreicht, zeigt das Beispiel der divergenten harmonischen Reihe $\sum_{n=1}^{\infty} \frac{1}{n}$. Mit $a_n := 1/n$ gilt zwar

$$\left| \frac{a_{n+1}}{a_n} \right| = \frac{n}{n+1} < 1 \quad \text{für alle } n \geq 1,$$

wegen $\lim_{n \rightarrow \infty} \frac{n}{n+1} = 1$ gibt es jedoch *kein* $\theta < 1$ mit

$$\left| \frac{a_{n+1}}{a_n} \right| \leq \theta \quad \text{für alle } n \geq n_0.$$

Das Quotienten-Kriterium ist also nicht anwendbar.

(7.8) Für $a_n := 1/n^2$ erhalten wir die Reihe $\sum_{n=1}^{\infty} a_n = \sum_{n=1}^{\infty} \frac{1}{n^2}$.

Sie konvergiert, wie wir bereits wissen. Auch hier ist

$$\left| \frac{a_{n+1}}{a_n} \right| = \frac{n^2}{(n+1)^2} < 1 \quad \text{für alle } n \geq 1,$$

es gibt aber kein $\theta < 1$ mit

$$\left| \frac{a_{n+1}}{a_n} \right| \leq \theta \quad \text{für alle } n \geq n_0.$$

Das Quotienten-Kriterium ist also nicht anwendbar, obwohl die Reihe konvergiert. Das bedeutet, dass das Quotienten-Kriterium nur eine hinreichende, jedoch nicht notwendige Bedingung für die Konvergenz ist.

Umordnung von Reihen

Sei $\sum_{n=0}^{\infty} a_n$ eine Reihe und $\tau: \mathbb{N} \rightarrow \mathbb{N}$ eine bijektive Abbildung. Dann nennt man $\sum_{n=0}^{\infty} a_{\tau(n)}$ eine *Umordnung* der gegebenen Reihe. Sie besteht aus denselben Summanden, nur in einer anderen Reihenfolge. Anders als bei endlichen Summen ist es bei konvergenten unendlichen Reihen nicht ohne weiteres klar, dass sie nach Umordnung wieder konvergent mit demselben Grenzwert sind. Für absolut konvergente Reihen ist dies jedoch richtig.

Satz 8 (Umordnungssatz). *Sei $\sum_{n=0}^{\infty} a_n$ eine absolut konvergente Reihe. Dann konvergiert auch jede Umordnung dieser Reihe absolut gegen denselben Grenzwert.*

Beweis. Sei $A := \sum_{n=0}^{\infty} a_n$ und $\tau: \mathbb{N} \rightarrow \mathbb{N}$ eine bijektive Abbildung. Wir müssen zeigen

$$\lim_{m \rightarrow \infty} \sum_{k=0}^m a_{\tau(k)} = A.$$

Sei $\varepsilon > 0$ vorgegeben. Dann gibt es wegen der Konvergenz von $\sum_{k=0}^{\infty} |a_k|$ ein $n_0 \in \mathbb{N}$, so dass

$$\sum_{k=n_0}^{\infty} |a_k| < \frac{\varepsilon}{2}.$$

Daraus folgt

$$\left| A - \sum_{k=0}^{n_0-1} a_k \right| = \left| \sum_{k=n_0}^{\infty} a_k \right| \leq \sum_{k=n_0}^{\infty} |a_k| < \frac{\varepsilon}{2}.$$

Sei N so groß gewählt, dass

$$\{\tau(0), \tau(1), \dots, \tau(N)\} \supset \{0, 1, 2, \dots, n_0 - 1\}.$$

Dann gilt für alle $m \geq N$

$$\left| \sum_{k=0}^m a_{\tau(k)} - A \right| \leq \left| \sum_{k=0}^m a_{\tau(k)} - \sum_{k=0}^{n_0-1} a_k \right| + \left| \sum_{k=0}^{n_0-1} a_k - A \right| \leq \sum_{k=n_0}^{\infty} |a_k| + \frac{\varepsilon}{2} < \varepsilon,$$

die umgeordnete Reihe konvergiert also gegen denselben Grenzwert wie die Ausgangsreihe. Dass die umgeordnete Reihe wieder absolut konvergiert, folgt aus der Anwendung des gerade Bewiesenen auf die Reihe $\sum_{n=0}^{\infty} |a_n|$.

(7.9) Wir zeigen an einem Beispiel, dass der Satz 8 falsch wird, wenn man nicht verlangt, dass die Reihe absolut konvergiert. Dazu verwenden wir die

nach (7.3) konvergente alternierende harmonische Reihe

$$\sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{n} = 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} \pm \dots$$

Behauptung. Es gibt eine Umordnung mit $\sum_{n=1}^{\infty} \frac{(-1)^{\tau(n)-1}}{\tau(n)} = \infty$.

Beweis. Wir betrachten die Glieder ungerader Ordnung der gegebenen Reihe von $\frac{1}{2^{n+1}}$ bis $\frac{1}{2^{n+1}-1}$. Für jedes $n \geq 1$ gilt

$$\frac{1}{2^{n+1}} + \frac{1}{2^{n+3}} + \dots + \frac{1}{2^{n+1}-1} > 2^{n-1} \cdot \frac{1}{2^{n+1}} = \frac{1}{4}.$$

Deshalb divergiert folgende Reihen-Umordnung bestimmt gegen $+\infty$:

$$\begin{aligned} & 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} \\ & \quad + \left(\frac{1}{5} + \frac{1}{7} \right) - \frac{1}{6} \\ & \quad + \left(\frac{1}{9} + \frac{1}{11} + \frac{1}{13} + \frac{1}{15} \right) - \frac{1}{8} \\ & \quad + \dots \\ & \quad + \left(\frac{1}{2^n + 1} + \frac{1}{2^n + 3} + \dots + \frac{1}{2^{n+1} - 1} \right) - \frac{1}{2n+2} \\ & \quad + \dots \end{aligned}$$

Man beachte, dass in der Umordnung alle mit Minuszeichen behafteten Glieder gerader Ordnung einmal an die Reihe kommen, aber mit immer größerer Verzögerung gegenüber den positiven Gliedern ungerader Ordnung. Deshalb können die Partialsummen über alle Grenzen wachsen.

Dies Gegenbeispiel zeigt also, dass für nicht absolut konvergente unendliche Summen das Kommutativgesetz nicht gilt. Siehe dazu auch Aufgabe 7.13.

AUFGABEN

7.1. Man untersuche die folgenden Reihen auf Konvergenz oder Divergenz:

$$\sum_{n=1}^{\infty} \frac{n!}{n^n}, \quad \sum_{n=0}^{\infty} \frac{n^4}{3^n}, \quad \sum_{n=0}^{\infty} \frac{n+4}{n^2-3n+1}, \quad \sum_{n=1}^{\infty} \frac{(n+1)^{n-1}}{(-n)^n}.$$

7.2. Sei $(a_n)_{n \geq 1}$ eine Folge reeller Zahlen mit $|a_n| \leq M$ für alle $n \geq 1$. Man zeige:

a) Für jedes $x \in \mathbb{R}$ mit $|x| < 1$ konvergiert die Reihe

$$f(x) := \sum_{n=1}^{\infty} a_n x^n.$$

b) Ist $a_1 \neq 0$, so gilt

$$f(x) \neq 0 \quad \text{für alle } x \in \mathbb{R} \text{ mit } 0 < |x| < \frac{|a_1|}{2M}.$$

7.3. (Reihenverdichtungs-Kriterium) Es sei $(a_n)_{n \in \mathbb{N}}$ eine Folge nicht-negativer reeller Zahlen mit $a_0 \geq a_1 \geq a_2 \geq \dots \geq a_k \geq a_{k+1} \geq \dots$.

Man beweise:

Die Reihe $\sum_{n=0}^{\infty} a_n$ konvergiert genau dann, wenn $\sum_{n=0}^{\infty} 2^n a_{2^n}$ konvergiert.

7.4. Man zeige: Für jede natürliche Zahl $k \geq 2$ konvergiert die Reihe

$$\sum_{n=1}^{\infty} \frac{1}{n^{\frac{1}{k}} n}.$$

7.5. (Wurzelkriterium) Sei $\sum_{n=0}^{\infty} a_n$ eine unendliche Reihe. Es gebe ein θ mit $0 < \theta < 1$ und ein $n_0 > 0$, so dass

$$\sqrt[n]{|a_n|} \leq \theta \quad \text{für alle } n \geq n_0.$$

a) Man beweise, dass die Reihe absolut konvergiert.

b) Man zeige, dass die Bedingung

$$\sqrt[n]{|a_n|} < 1 \quad \text{für alle } n \geq n_0.$$

nicht hinreichend für die Konvergenz der Reihe $\sum a_n$ ist. Ist dies eine notwendige Bedingung?

7.6. Es sei $\sum_{n=0}^{\infty} a_n$ eine unendliche Reihe, deren absolute Konvergenz mit dem Quotienten-Kriterium bewiesen werden kann.

Man zeige: Dann ist auch das Wurzel-Kriterium anwendbar.

Gilt auch die Umkehrung?

7.7. (Raabesches Konvergenz-Kriterium) Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge positiver reeller Zahlen. Man beweise

a) Es gebe eine reelle Zahl $\alpha > 1$ und ein $n_0 \geq 2$, so dass

$$\frac{a_n}{a_{n-1}} \leq 1 - \frac{\alpha}{n} \quad \text{für alle } n \geq n_0.$$

Dann ist die Reihe $\sum_{n=0}^{\infty} a_n$ konvergent.

b) Falls

$$\frac{a_n}{a_{n-1}} \geq 1 - \frac{1}{n} \quad \text{für alle } n \geq n_0,$$

so divergiert die Reihe $\sum_{n=0}^{\infty} a_n$.

7.8. Man beweise mithilfe des Raabeschen Konvergenz-Kriteriums: Die Reihe

$$\sum_{n=1}^{\infty} \binom{1/2}{n}$$

konvergiert absolut. (Zur Definition von $\binom{1/2}{n}$ vgl. Aufgabe 1.11.)

7.9. Es bezeichne

$$M_1 = \{2, 3, 4, \dots, 8, 9, 20, 22, 23, \dots, 29, 30, 32, \dots, 39, 40, 42, \dots\}$$

die Menge aller positiven ganzen Zahlen, in deren Dezimaldarstellung die Ziffer 1 nicht vorkommt. Man zeige

$$\sum_{n \in M_1} \frac{1}{n} < \infty.$$

7.10. Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge reeller Zahlen mit $\lim_{n \rightarrow \infty} a_n = 0$. Die Folge $(A_k)_{k \in \mathbb{N}}$ werde definiert durch

$$A_0 := \frac{1}{2}a_0,$$

$$A_k := \frac{1}{2}a_{2k-2} + a_{2k-1} + \frac{1}{2}a_{2k} \quad \text{für } k \geq 1.$$

Man beweise: Konvergiert eine der beiden Reihen

$$\sum_{n=0}^{\infty} a_n, \quad \sum_{k=0}^{\infty} A_k,$$

so konvergiert auch die zweite gegen denselben Grenzwert.

7.11. Unter Benutzung der Summe der Leibniz'schen Reihe (7.4) beweise man

$$\pi = 2 + \sum_{k=1}^{\infty} \frac{16}{(4k-3)(16k^2-1)}.$$

Man vergleiche die Konvergenz-Geschwindigkeit der Leibniz'schen Reihe und der obigen Reihe. (Die Reihe eignet sich gut zu kleinen Programmier-Experimenten!)

7.12. Die bijektive Abbildung $\tau : \mathbb{N} \longrightarrow \mathbb{N}$ sei eine beschränkte Umordnung, d.h. es gebe ein $d \in \mathbb{N}$, so dass

$$|\tau(n) - n| \leq d \quad \text{für alle } n \in \mathbb{N}.$$

Man beweise: Eine Reihe $\sum_{n=0}^{\infty} a_n$ konvergiert genau dann, wenn die Reihe $\sum_{n=0}^{\infty} a_{\tau(n)}$ konvergiert.

7.13. Sei $\sum a_n$ eine konvergente, aber nicht absolut konvergente Reihe reeller Zahlen. Man beweise:

- Zu beliebig vorgegebenem $c \in \mathbb{R}$ gibt es eine Umordnung $\sum a_{\tau(n)}$, die gegen c konvergiert.
- Es gibt Umordnungen, so dass $\sum a_{\tau(n)}$ bestimmt gegen $+\infty$ bzw. $-\infty$ divergiert.
- Es gibt Umordnungen, so dass $\sum a_{\tau(n)}$ weder konvergiert noch bestimmt gegen $\pm\infty$ divergiert.

§ 8 Die Exponentialreihe

Wir behandeln jetzt die Exponentialreihe, die neben der geometrischen Reihe die wichtigste Reihe in der Analysis ist. Die Funktionalgleichung der Exponentialfunktion beweisen wir mithilfe eines allgemeinen Satzes über das sog. Cauchy-Produkt von Reihen.

Satz 1. Für jedes $x \in \mathbb{R}$ ist die Exponentialreihe

$$\exp(x) := \sum_{n=0}^{\infty} \frac{x^n}{n!}$$

absolut konvergent.

Beweis. Die Behauptung folgt aus dem Quotienten-Kriterium (§7, Satz 7). Mit $a_n := x^n/n!$ gilt für alle $x \neq 0$ und $n \geq 2|x|$

$$\left| \frac{a_{n+1}}{a_n} \right| = \left| \frac{x^{n+1}}{(n+1)!} \cdot \frac{n!}{x^n} \right| = \frac{|x|}{n+1} \leq \frac{1}{2}, \quad \text{q.e.d.}$$

Mit der Exponentialreihe definiert man die berühmte *Eulersche Zahl*

$$e := \exp(1) = \sum_{n=0}^{\infty} \frac{1}{n!} = 1 + 1 + \frac{1}{2} + \frac{1}{3!} + \dots = 2.7182818\dots$$

Bemerkung. Eine äquivalente Definition $e = \lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n$ werden wir später in (15.12) kennenlernen.

Satz 2 (Abschätzung des Restglieds). Es gilt

$$\exp(x) = \sum_{n=0}^N \frac{x^n}{n!} + R_{N+1}(x),$$

wobei

$$|R_{N+1}(x)| \leq 2 \frac{|x|^{N+1}}{(N+1)!} \quad \text{für alle } x \text{ mit } |x| \leq 1 + \frac{1}{2}N.$$

Bei Abbruch der Reihe ist also der Fehler in dem angegebenen x -Bereich dem Betrage nach höchstens zweimal so groß wie das erste nicht berücksichtigte Glied.

Beweis. Wir schätzen den Rest $R_{N+1}(x) = \sum_{n=N+1}^{\infty} x^n/n!$ mittels der geometrischen Reihe ab. Es ist

$$\begin{aligned} |R_{N+1}(x)| &\leq \sum_{n=N+1}^{\infty} \frac{|x|^n}{n!} \\ &= \frac{|x|^{N+1}}{(N+1)!} \left\{ 1 + \frac{|x|}{N+2} + \frac{|x|^2}{(N+2)(N+3)} + \dots \right\} \\ &\leq \frac{|x|^{N+1}}{(N+1)!} \left\{ 1 + \frac{|x|}{N+2} + \left(\frac{|x|}{N+2}\right)^2 + \left(\frac{|x|}{N+2}\right)^3 + \dots \right\}. \end{aligned}$$

Für $|x| \leq 1 + \frac{1}{2}N$ ist der Ausdruck innerhalb der geschweiften Klammer $\leq 1 + \frac{1}{2} + \frac{1}{4} + \frac{1}{8} + \dots = 2$, woraus die Behauptung folgt.

Numerische Berechnung von e

Wir benutzen die Fehler-Abschätzung, um $e = \exp(1)$ mit hoher Genauigkeit zu berechnen. Aus Satz 2 folgt für $k \geq 1$, dass $R_{k+1}(1) \leq 1/k!$, d.h. bei Abbruch der Reihe für e ist der Fehler höchstens so groß wie das letzte berücksichtigte Glied. Die Partialsummen $S_k := \sum_{v=0}^k 1/v!$ kann man rekursiv durch

$$S_0 := u_0 := 1, \quad u_k := \frac{u_{k-1}}{k}, \quad S_k := S_{k-1} + u_k, \quad (k > 0),$$

berechnen. Um eine vorgegebene Fehlerschranke $\varepsilon > 0$ zu unterschreiten, braucht man nur solange zu rechnen, bis $u_k < \varepsilon$ wird. Wir schreiben eine ARIBAS-Funktion `euler(n)`, die e auf n Dezimalstellen mit einem Fehler $\leq 10^{-n}$ ausrechnet. Dabei verwenden wir Ganzzahl-Arithmetik und multipli-

```
function euler(n: integer): integer;
var
  S, u, k: integer;
begin
  S := u := 10**(n+5);
  k := 0;
  while u > 0 do
    k := k+1;
    u := u div k;
    S := S+u;
  end;
  writeln("Euler number calculated in ",k," steps");
  return (S div 10**5);
end.
```

zieren alle Größen mit 10^n . Dann braucht nur bis auf $\varepsilon = 1$ genau gerechnet zu werden. Zur Berücksichtigung von Rundungsfehlern rechnen wir noch mit 5 Stellen mehr. In diesem Code ist `u div k` (wie in PASCAL) die Integer-Division, d.h. es wird die ganze Zahl $\lfloor u/k \rfloor$ berechnet, die vom exakten Ergebnis u/k um weniger als 1 abweicht. Da u ganzzahlig ist, wird die while-Schleife abgebrochen, sobald $u < 1$ ist, und der gesamte akkumulierte Rundungsfehler ist höchstens gleich der Anzahl der Schleifen-Durchgänge. Solange diese kleiner als 10^5 bleibt, wird die angestrebte Genauigkeit erreicht. Testen wir die Funktion mit $n = 100$, ergibt sich

```
==> euler(100) .
Euler number calculated in 73 steps
-: 2_71828_18284_59045_23536_02874_71352_66249_77572_
47093_69995_95749_66967_62772_40766_30353_54759_45713_
82178_52516_64274
```

Dies ist natürlich als $e = 2.71828\dots$ zu interpretieren. Hier wurde also in 73 Schritten e auf 100 Dezimalstellen genau berechnet. Wir ersparen uns Tests mit höherer Stellenzahl (etwa $n = 1000$ oder $n = 10000$), die der Leser leicht selbst durchführen kann.

Cauchy-Produkt von Reihen

Zum Beweis der Funktionalgleichung der Exponentialfunktion benützen wir folgenden allgemeinen Satz über das Produkt von unendlichen Reihen.

Satz 3 (Cauchy-Produkt von Reihen). *Es seien $\sum_{n=0}^{\infty} a_n$ und $\sum_{n=0}^{\infty} b_n$ absolut konvergente Reihen. Für $n \in \mathbb{N}$ werde definiert*

$$c_n := \sum_{k=0}^n a_k b_{n-k} = a_0 b_n + a_1 b_{n-1} + \dots + a_n b_0.$$

Dann ist auch die Reihe $\sum_{n=0}^{\infty} c_n$ absolut konvergent mit

$$\sum_{n=0}^{\infty} c_n = \left(\sum_{n=0}^{\infty} a_n \right) \cdot \left(\sum_{n=0}^{\infty} b_n \right).$$

Beweis. Die Definition des Koeffizienten c_n lässt sich auch so schreiben:

$$c_n = \sum \{ a_k b_\ell : k + \ell = n \}.$$

Es wird dabei über alle Indexpaare (k, ℓ) summiert, die in $\mathbb{N} \times \mathbb{N}$ auf der Diagonalen $k + \ell = n$ liegen. Deshalb gilt für die Partialsumme

$$C_N := \sum_{n=0}^N c_n = \sum \{ a_k b_\ell : (k, \ell) \in \Delta_N \},$$

wobei Δ_N das wie folgt definierte Dreieck in $\mathbb{N} \times \mathbb{N}$ ist:

$$\Delta_N := \{(k, \ell) \in \mathbb{N} \times \mathbb{N} : k + \ell \leq N\}, \quad \text{vgl. Bild 8.1.}$$

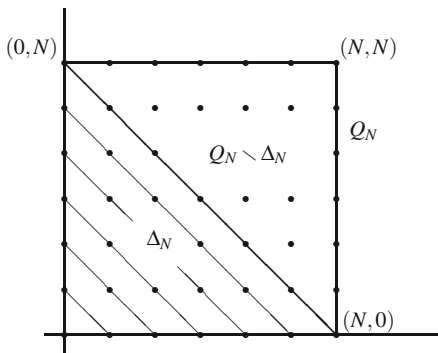


Bild 8.1

Multiplizieren wir die Partialsummen

$$A_N := \sum_{n=0}^N a_n \quad \text{und} \quad B_N := \sum_{n=0}^N b_n$$

aus, erhalten wir als Produkt

$$A_N B_N = \sum \{a_k b_\ell : (k, \ell) \in Q_N\},$$

wobei Q_N das Quadrat

$$Q_N := \{(k, \ell) \in \mathbb{N} \times \mathbb{N} : 0 \leq k \leq N, 0 \leq \ell \leq N\}$$

bezeichnet. Da $\Delta_N \subset Q_N$, können wir schreiben

$$A_N B_N - C_N = \sum \{a_k b_\ell : (k, \ell) \in Q_N \setminus \Delta_N\}.$$

Für die Partialsummen

$$A_N^* := \sum_{n=0}^N |a_n|, \quad B_N^* := \sum_{n=0}^N |b_n|$$

erhält man wie oben

$$A_N^* B_N^* = \sum \{|a_k| |b_\ell| : (k, \ell) \in Q_N\}.$$

Da $Q_{\lfloor N/2 \rfloor} \subset \Delta_N$, folgt $Q_N \setminus \Delta_N \subset Q_N \setminus Q_{\lfloor N/2 \rfloor}$, also

$$\begin{aligned}
 |A_N B_N - C_N| &\leq \sum \{|a_k| |b_\ell| : (k, \ell) \in Q_N \setminus Q_{[N/2]}\} \\
 &= A_N^* B_N^* - A_{[N/2]}^* B_{[N/2]}^*.
 \end{aligned}$$

Da die Folge $(A_N^* B_N^*)$ konvergiert, also eine Cauchy-Folge ist, strebt die letzte Differenz für $N \rightarrow \infty$ gegen 0, d.h.

$$\lim_{N \rightarrow \infty} C_N = \lim_{N \rightarrow \infty} A_N B_N = \lim_{N \rightarrow \infty} A_N \lim_{N \rightarrow \infty} B_N.$$

Damit ist gezeigt, dass $\sum c_n$ konvergiert und die im Satz behauptete Formel über das Cauchy-Produkt gilt. Es ist noch die absolute Konvergenz von $\sum c_n$ zu beweisen. Wegen

$$|c_n| \leq \sum_{k=0}^n |a_k| |b_{n-k}|$$

ergibt sich dies durch Anwendung des bisher Bewiesenen auf die Reihen $\sum |a_n|$ und $\sum |b_n|$.

Bemerkung. Die Voraussetzung der absoluten Konvergenz ist wesentlich für die Gültigkeit von Satz 3, vgl. Aufgabe 8.2.

Satz 4 (Funktionalgleichung der Exponentialfunktion). *Für alle $x, y \in \mathbb{R}$ gilt*

$$\exp(x+y) = \exp(x) \exp(y).$$

Bemerkung. Diese Funktionalgleichung heißt auch *Additions-Theorem* der Exponentialfunktion.

Beweis. Wir bilden das Cauchy-Produkt der absolut konvergenten Reihen $\exp(x) = \sum x^n/n!$ und $\exp(y) = \sum y^n/n!$. Für den n -ten Koeffizienten der Produktreihe ergibt sich mit dem binomischen Lehrsatz

$$c_n = \sum_{k=0}^n \frac{x^k}{k!} \cdot \frac{y^{n-k}}{(n-k)!} = \frac{1}{n!} \sum_{k=0}^n \binom{n}{k} x^k y^{n-k} = \frac{1}{n!} (x+y)^n.$$

Also folgt $\exp(x) \exp(y) = \sum \frac{1}{n!} (x+y)^n = \exp(x+y)$, q.e.d.

Corollar. a) *Für alle $x \in \mathbb{R}$ gilt $\exp(x) > 0$.*

b) *Für alle $x \in \mathbb{R}$ gilt $\exp(-x) = \frac{1}{\exp(x)}$.*

c) *Für jede ganze Zahl $n \in \mathbb{Z}$ ist $\exp(n) = e^n$.*

Beweis. b) Aufgrund der Funktionalgleichung ist

$$\exp(x) \exp(-x) = \exp(x-x) = \exp(0) = 1,$$

also insbesondere $\exp(x) \neq 0$ und $\exp(-x) = \exp(x)^{-1}$.

a) Für $x \geq 0$ sieht man an der Reihendarstellung, dass

$$\exp(x) = 1 + x + \frac{x^2}{2} + \dots \geq 1 > 0.$$

Ist $x < 0$, so folgt $-x > 0$ also $\exp(-x) > 0$ und damit

$$\exp(x) = \exp(-x)^{-1} > 0.$$

c) Wir zeigen zunächst mit vollständiger Induktion, dass für alle $n \in \mathbb{N}$ gilt $\exp(n) = e^n$.

Induktionsanfang $n = 0$. Es ist $\exp(0) = 1 = e^0$.

Induktionsschritt $n \rightarrow n+1$. Mit der Funktionalgleichung und Induktionsvoraussetzung erhält man

$$\exp(n+1) = \exp(n) \exp(1) = e^n e = e^{n+1}.$$

Damit ist $\exp(n) = e^n$ für $n \geq 0$ bewiesen. Mittels b) ergibt sich daraus

$$\exp(-n) = \frac{1}{\exp(n)} = \frac{1}{e^n} = e^{-n} \quad \text{für alle } n \in \mathbb{N}.$$

Somit gilt $\exp(n) = e^n$ für alle ganzen Zahlen n .

Bemerkung. Die Formel c) des Corollars motiviert die Bezeichnung Exponentialfunktion. Man kann sagen, dass $\exp(x)$ die Potenzen e^n , $n \in \mathbb{Z}$, interpoliert und so auf nicht-ganze Exponenten ausdehnt. Man schreibt deshalb auch suggestiv e^x für $\exp(x)$. Die Formel zeigt auch, dass es genügt, die Werte der Exponentialfunktion im Bereich $-\frac{1}{2} \leq x \leq \frac{1}{2}$ zu kennen, um sie für alle x zu kennen. Denn jedes $x \in \mathbb{R}$ lässt sich schreiben als $x = n + \xi$ mit $n \in \mathbb{Z}$ und $|\xi| \leq \frac{1}{2}$ und es gilt dann

$$\exp(x) = \exp(n + \xi) = e^n \exp(\xi).$$

Da $|\xi|$ klein ist, konvergiert die Exponentialreihe für $\exp(\xi)$ besonders schnell.

AUFGABEN

8.1. a) Sei $x \geq 1$ eine reelle Zahl. Man zeige, dass die Reihe

$$s(x) := \sum_{n=0}^{\infty} \binom{x}{n}$$

absolut konvergiert. (Die Zahlen $\binom{x}{n}$ wurden in Aufgabe 1.11 definiert.)

b) Man beweise die Funktionalgleichung

$$s(x+y) = s(x)s(y) \quad \text{für alle } x, y \geq 1.$$

c) Man berechne $s(n + \frac{1}{2})$ für alle natürlichen Zahlen $n \geq 1$.

8.2. Für $n \in \mathbb{N}$ sei

$$a_n := b_n := \frac{(-1)^n}{\sqrt{n+1}} \quad \text{und} \quad c_n := \sum_{k=0}^n a_{n-k} b_k.$$

Man zeige, dass die Reihen $\sum_{n=0}^{\infty} a_n$ und $\sum_{n=0}^{\infty} b_n$ konvergieren, ihr Cauchy-Produkt $\sum_{n=0}^{\infty} c_n$ aber nicht konvergiert.

8.3. Sei $A(n)$ die Anzahl aller Paare $(k, \ell) \in \mathbb{N} \times \mathbb{N}$ mit

$$n = k^2 + \ell^2.$$

Man beweise: Für alle x mit $|x| < 1$ gilt

$$\left(\sum_{n=0}^{\infty} x^{n^2} \right)^2 = \sum_{n=0}^{\infty} A(n) x^n.$$

8.4. Sei $M = \{1, 2, 4, 5, 8, 10, 16, 20, 25, \dots\}$ die Menge aller natürlichen Zahlen ≥ 1 , die durch keine Primzahl $\neq 2, 5$ teilbar sind. Man betrachte die zu M gehörige Teilreihe der harmonischen Reihe und beweise

$$\sum_{n \in M} \frac{1}{n} = \frac{5}{2}.$$

Anleitung. Man bilde das Produkt der geometrischen Reihen $\sum 2^{-n}$ und $\sum 5^{-n}$.

8.5. (Verallgemeinerung von Aufgabe 8.4.) Sei $\mathcal{P}$ eine endliche Menge von Primzahlen und $\mathcal{N}(\mathcal{P})$ die Menge aller natürlichen Zahlen ≥ 1 , in deren Primfaktor-Zerlegung höchstens Primzahlen aus $\mathcal{P}$ vorkommen (Existenz und Eindeutigkeit der Primfaktor-Zerlegung sei vorausgesetzt.) Man beweise, dass

$$\sum_{n \in \mathcal{N}(\mathcal{P})} \frac{1}{n} = \prod_{p \in \mathcal{P}} \left(1 - \frac{1}{p}\right)^{-1} < \infty.$$

Bemerkung. Ist $\mathcal{P}$ die Menge aller Primzahlen, so besteht $\mathcal{N}(\mathcal{P})$ aus allen natürlichen Zahlen ≥ 1 . Daraus kann man nach Euler folgern, dass es unendlich viele Primzahlen gibt. Gäbe es nur endlich viele, würde die harmonische Reihe konvergieren.

§ 9 Punktmengen

In diesem Paragraphen behandeln wir die Begriffe Abzählbarkeit und Überabzählbarkeit und beweisen insbesondere, dass die Menge aller reellen Zahlen nicht abzählbar ist. Weiter beschäftigen wir uns mit dem Supremum und Infimum von Mengen reeller Zahlen und definieren den Limes superior und Limes inferior von Folgen.

Bezeichnungen. Wir verwenden folgende Bezeichnungen für Intervalle auf der Zahlengeraden $\mathbb{R}$.

a) *Abgeschlossene Intervalle.* Seien $a, b \in \mathbb{R}$, $a \leq b$. Dann setzt man

$$[a, b] := \{x \in \mathbb{R} : a \leq x \leq b\}.$$

Für $a = b$ besteht $[a, b]$ nur aus einem Punkt.

b) *Offene Intervalle.* Seien $a, b \in \mathbb{R}$, $a < b$. Man setzt

$$]a, b[:= \{x \in \mathbb{R} : a < x < b\}.$$

c) *Halboffene Intervalle.* Für $a, b \in \mathbb{R}$, $a < b$ sei

$$\begin{aligned} [a, b[&:= \{x \in \mathbb{R} : a \leq x < b\}, \\]a, b] &:= \{x \in \mathbb{R} : a < x \leq b\}. \end{aligned}$$

d) *Uneigentliche Intervalle.* Sei $a \in \mathbb{R}$. Man definiert

$$\begin{aligned} [a, +\infty[&:= \{x \in \mathbb{R} : x \geq a\}, \\]a, +\infty[&:= \{x \in \mathbb{R} : x > a\}, \\]-\infty, a] &:= \{x \in \mathbb{R} : x \leq a\}, \\]-\infty, a[&:= \{x \in \mathbb{R} : x < a\}. \end{aligned}$$

Weitere Bezeichnungen:

$$\begin{aligned} \mathbb{R}_+ &:= \{x \in \mathbb{R} : x \geq 0\}, \\ \mathbb{R}^* &:= \{x \in \mathbb{R} : x \neq 0\}, \\ \mathbb{R}_+^* &:= \mathbb{R}_+ \cap \mathbb{R}^* = \{x \in \mathbb{R} : x > 0\} =]0, +\infty[. \end{aligned}$$

Die reelle Zahlengerade $\mathbb{R}$ wird manchmal auch mit $]-\infty, +\infty[$ bezeichnet.

Mengen und Folgen

Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge reeller Zahlen. Dann heißt die Menge

$$M := \{a_n : n \in \mathbb{N}\}$$

die der Folge $(a_n)_{n \in \mathbb{N}}$ *unterliegende Menge*. Verschiedene Folgen können dieselbe unterliegende Menge haben, wie folgendes Beispiel zeigt:

Die Folgen $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ mit $a_n = (-1)^n$ und $b_n = (-1)^{n+1}$ sind verschieden, sie haben jedoch beide dieselbe unterliegende Menge, nämlich $\{-1, +1\}$.

Abzählbarkeit, Überabzählbarkeit

Definition. Eine nichtleere Menge A heißt *abzählbar*, wenn es eine surjektive Abbildung $\mathbb{N} \rightarrow A$ gibt, d.h. wenn eine Folge $(x_n)_{n \in \mathbb{N}}$ existiert, so dass $A = \{x_n : n \in \mathbb{N}\}$. Die leere Menge wird ebenfalls als abzählbar definiert. Eine nichtleere Menge heißt *überabzählbar*, wenn sie nicht abzählbar ist.

Bemerkung. Jede endliche Menge $A = \{a_0, \dots, a_n\}$ ist abzählbar. Eine surjektive Abbildung $\tau: \mathbb{N} \rightarrow A$ kann man wie folgt definieren: Man setze $\tau(k) := a_k$ für $0 \leq k \leq n$ und $\tau(k) := a_n$ für $k > n$. Eine nicht-endliche abzählbare Menge nennt man abzählbar unendlich. Man kann sich leicht überlegen, dass es zu jeder abzählbar unendlichen Menge M sogar eine bijektive Abbildung $\mathbb{N} \rightarrow M$ gibt.

Beispiele

(9.1) Die Menge $\mathbb{N}$ aller natürlichen Zahlen ist abzählbar, denn die identische Abbildung $\mathbb{N} \rightarrow \mathbb{N}$ ist surjektiv. Jede Teilmenge $M \subset \mathbb{N}$ ist entweder endlich oder abzählbar unendlich. Ist $M \subset \mathbb{N}$ eine nicht-endliche Teilmenge, so erhält man eine bijektive Abbildung $\sigma: M \rightarrow \mathbb{N}$ z.B. so: Für $n \in M$ sei $\sigma(n) = k$, wobei k die Anzahl der Elemente von

$$M_n := \{x \in M : x < n\}$$

ist.

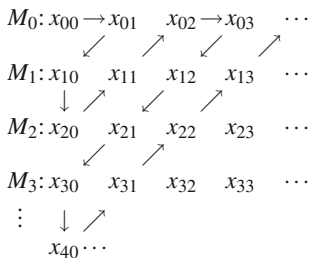
(9.2) Die Menge $\mathbb{Z}$ aller ganzen Zahlen ist abzählbar. Denn $\mathbb{Z}$ ist unterliegende Menge der Folge

$$(x_0, x_1, x_2, \dots) := (0, +1, -1, +2, -2, \dots).$$

(Es ist $x_0 = 0$ und $x_{2k-1} = k$, $x_{2k} = -k$ für $k \geq 1$.)

Satz 1. Die Vereinigung abzählbar vieler abzählbarer Mengen M_n , $n \in \mathbb{N}$, ist wieder abzählbar.

Beweis. Sei $M_n = \{x_{nm} : m \in \mathbb{N}\}$. Wir schreiben die Vereinigungsmenge $\bigcup_{n \in \mathbb{N}} M_n$ in einem quadratisch unendlichen Schema an.



Die durch die Pfeile angedeutete Abzählungsvorschrift liefert eine Folge

$$(y_0, y_1, y_2, \dots) := (x_{00}, x_{01}, x_{10}, x_{20}, x_{11}, x_{02}, x_{03}, x_{12}, \dots),$$

für die $\bigcup_{n \in \mathbb{N}} M_n = \{y_n : n \in \mathbb{N}\}$, q.e.d.

Corollar 1. Die Menge $\mathbb{Q}$ der rationalen Zahlen ist abzählbar.

Beweis. Da die Menge $\mathbb{Z}$ aller ganzen Zahlen abzählbar ist, ist für jede feste natürliche Zahl $k \geq 1$ die Menge

$$A_k := \left\{ \frac{n}{k} : n \in \mathbb{Z} \right\}$$

abzählbar. Da $\mathbb{Q} = \bigcup_{k \geq 1} A_k$, ist nach Satz 1 auch $\mathbb{Q}$ abzählbar.

Aus Satz 1 lassen sich weitere interessante Folgerungen ziehen. Betrachten wir etwa die Menge aller in einer gewissen Programmiersprache, z.B. in PASCAL, geschriebenen Programme. Jedes einzelne solche Programm kann man darstellen als eine endliche Folge von Bytes. (Natürlich ist nicht jede endliche Byte-Folge ein gültiges PASCAL-Programm.) Daraus folgt, dass die Menge P_n aller PASCAL-Programme, die eine Länge von n Bytes haben, endlich, also auch abzählbar ist. Daher ist die Vereinigung $\bigcup_{n \geq 1} P_n$ ebenfalls abzählbar. Damit haben wir bewiesen:

Corollar 2. *Die Menge aller möglichen PASCAL-Programme ist abzählbar.*

Nennt man eine reelle Zahl x *berechenbar*, wenn es ein Programm gibt, das bei Eingabe einer natürlichen Zahl n die Zahl x mit einem Fehler $\leq 2^{-n}$ berechnet (ähnlich wie wir es in §8 für die Zahl e durchgeführt haben), so folgt, dass es nur abzählbar viele berechenbare reelle Zahlen gibt, was im Kontrast zu dem nachfolgend bewiesenen Satz steht, dass die Menge aller reellen Zahlen überabzählbar ist. Nicht alle reellen Zahlen sind also berechenbar.

Bemerkung. Natürlich kann man auf einem konkreten Computer wegen der Endlichkeit des Speicherplatzes i.Allg. eine reelle Zahl nicht mit beliebiger Genauigkeit berechnen. Deshalb legt man in der theoretischen Informatik das Modell der sog. *Turing-Maschine* zugrunde, die zwar nur endlich viele Zustände hat und von einem endlichen Programm gesteuert wird, aber für die Ein- und Ausgabe ein nach zwei Richtungen unendliches Band zur Verfügung hat. Zu jedem Zeitpunkt sind aber nur endlich viele Zellen des Bandes beschrieben. (Eine Definition der Turing-Maschinen findet man in fast allen Lehrbüchern der Theoretischen Informatik, z.B. in [HU]. Der Leserin, die sich für die logischen Aspekte der Berechenbarkeit reeller Zahlen interessiert, sei das Buch [Br] empfohlen.)

Satz 2. *Die Menge $\mathbb{R}$ aller reellen Zahlen ist überabzählbar.*

Beweis. Wir zeigen, dass sogar das Intervall $I :=]0, 1[$ nicht abzählbar ist. Dazu genügt es offenbar, folgendes zu beweisen: Zu jeder Folge $(x_n)_{n \geq 1}$ von Zahlen $x_n \in I$ gibt es eine Zahl $z \in I$ mit $z \neq x_n$ für alle $n \geq 1$. Zum Beweis verwenden wir das sog. Cantorsche Diagonalverfahren. Die Dezimalbruch-Entwicklungen der Zahlen x_n seien

$$x_1 = 0.a_{11}a_{12}a_{13}\dots$$

$$x_2 = 0.a_{21}a_{22}a_{23}\dots$$

$$x_3 = 0.a_{31}a_{32}a_{33}\dots$$

$$\vdots$$

Wir definieren die Zahl $z \in I$ durch die Dezimalbruch-Entwicklung

$$z = 0.c_1c_2c_3\dots,$$

wobei

$$c_n := \begin{cases} a_{nn} + 2, & \text{falls } a_{nn} < 5, \\ a_{nn} - 2, & \text{falls } a_{nn} \geq 5. \end{cases}$$

Es gilt also $|c_n - a_{nn}| = 2$ für alle $n \geq 1$, woraus folgt $|z - x_n| \geq 10^{-n}$ für alle n , d.h. z ist verschieden von allen Gliedern der Folge (x_n) , q.e.d.

Bemerkung. Der Beweis zeigt, dass jedes nichtleere offene Intervall $]a, b[$ überabzählbar ist. Denn durch die Zuordnung

$$x \mapsto \frac{x-a}{b-a}$$

wird $]a, b[$ bijektiv auf $]0, 1[$ abgebildet.

Corollar. Die Menge der irrationalen Zahlen ist überabzählbar.

Beweis. Angenommen, die Menge $\mathbb{R} \setminus \mathbb{Q}$ der irrationalen Zahlen wäre abzählbar. Da $\mathbb{Q}$ abzählbar ist, wäre nach Satz 1 auch $\mathbb{R} = (\mathbb{R} \setminus \mathbb{Q}) \cup \mathbb{Q}$ abzählbar, Widerspruch!

Bemerkung. Ebenso zeigt man, dass die Menge der nicht berechenbaren reellen Zahlen überabzählbar ist, es gibt also mehr nicht berechenbare als berechenbare reelle Zahlen. Zur Beruhigung des Lesers sei jedoch gesagt, dass alle interessanten reellen Zahlen berechenbar sind. (Natürlich ist es eine Frage des Geschmacks, welche Zahlen man als interessant betrachtet, und darüber lässt sich streiten ...)

Berührungspunkte, Häufungspunkte

Definition. Sei $A \subset \mathbb{R}$ eine Teilmenge der Zahlengeraden und $a \in \mathbb{R}$.

a) Der Punkt a heißt *Berührungspunkt* von A , falls in jeder ε -Umgebung von a

$$U_\varepsilon(a) :=]a - \varepsilon, a + \varepsilon[, \quad \varepsilon > 0,$$

mindestens ein Punkt von A liegt.

b) Der Punkt a heißt *Häufungspunkt* von A , falls in jeder ε -Umgebung von a unendlich viele Punkte von A liegen.

Bemerkungen.

1) Jeder Punkt $a \in A$ ist trivialerweise Berührungspunkt von A .

2) a ist genau dann Berührungspunkt von A , wenn es eine Folge (a_n) von Punkten $a_n \in A$ mit $\lim_{n \rightarrow \infty} a_n = a$ gibt. Ist nämlich die Bedingung a) der Definition erfüllt, so wähle man für jedes $n \geq 1$ einen Punkt $a_n \in U_{1/n}(a) \cap A$. Damit gilt offensichtlich $\lim_{n \rightarrow \infty} a_n = a$. Die Umkehrung ist klar.

3) a ist genau dann Häufungspunkt von A , wenn a Berührungspunkt von $A \setminus \{a\}$ ist. Dann gibt es nämlich eine Folge $(a_n)_{n \in \mathbb{N}}$ von Punkten $a_n \in A \setminus \{a\}$ mit $\lim_{n \rightarrow \infty} a_n = a$, und unter diesen Punkten müssen unendlich viele verschiedene sein.

4) Nach der Definition in § 5 ist ein Punkt $a \in \mathbb{R}$ Häufungspunkt einer Folge $(a_n)_{n \in \mathbb{N}}$, wenn es eine Teilfolge gibt, die gegen a konvergiert. Dies kann man so umformulieren: a ist genau dann Häufungspunkt der Folge (a_n) , wenn es zu jeder ε -Umgebung von a unendlich viele Indizes $n \in \mathbb{N}$ gibt, so dass $a_n \in U_\varepsilon(a)$. Daraus folgt jedoch nicht, dass a Häufungspunkt der der Folge (a_n) unterliegenden Menge $\{a_n : n \in \mathbb{N}\}$ sein muss, denn diese Menge könnte endlich sein (z.B. für die konstante Folge $a_n = a$ für alle n).

Beispiele

(9.3) Die beiden Randpunkte a, b des offenen Intervalls $I =]a, b[\subset \mathbb{R}$, ($a < b$), sind Berührungspunkte und Häufungspunkte von I , außerdem alle Punkte von I selbst.

(9.4) Die Menge $A := \{\frac{1}{n} : n \in \mathbb{N}, n \geq 1\}$ hat 0 als einzigen Häufungspunkt, die Menge aller Berührungspunkte von A ist $A \cup \{0\}$.

(9.5) Jede reelle Zahl $x \in \mathbb{R}$ ist Häufungspunkt der Menge $\mathbb{Q}$ der rationalen Zahlen. Jedes $x \in \mathbb{R}$ ist auch Häufungspunkt der Menge $\mathbb{R} \setminus \mathbb{Q}$ der irrationalen Zahlen. Man drückt dies auch so aus: Sowohl die rationalen Zahlen als auch die irrationalen Zahlen liegen *dicht* in $\mathbb{R}$.

Supremum und Infimum von Punktmengen

Definition. Eine Teilmenge $A \subset \mathbb{R}$ heißt *nach oben* (bzw. *nach unten*) *beschränkt*, wenn es eine Konstante $K \in \mathbb{R}$ gibt, so dass

$$x \leq K \quad (\text{bzw. } x \geq K) \quad \text{für alle } x \in A.$$

Man nennt dann K obere (bzw. untere) *Schranke* von A . Die Menge A heißt *beschränkt*, wenn sie sowohl nach oben als auch nach unten beschränkt ist.

Bemerkungen. a) Eine Teilmenge $A \subset \mathbb{R}$ ist genau dann beschränkt, wenn es eine Konstante $M \geq 0$ gibt, so dass $|x| \leq M$ für alle $x \in A$.

b) Eine Folge ist genau dann nach oben (bzw. unten) beschränkt, wenn die ihr unterliegende Menge nach oben (bzw. unten) beschränkt ist (vgl. die Definition der Beschränktheit von Folgen in §4).

Definition. Sei A eine Teilmenge von $\mathbb{R}$. Eine Zahl $K \in \mathbb{R}$ heißt *Supremum* (bzw. *Infimum*) von A , falls K kleinste obere (bzw. größte untere) Schranke von A ist.

Dabei heißt K kleinste obere Schranke von A , falls gilt:

- i) K ist eine obere Schranke von A .
- ii) Ist K' eine weitere obere Schranke von A , so folgt $K \leq K'$.

Analog ist die größte untere Schranke von A definiert.

Es ist klar, dass die kleinste obere Schranke (bzw. größte untere Schranke) im Falle der Existenz eindeutig bestimmt ist. Man bezeichnet sie mit $\sup(A)$ bzw. $\inf(A)$.

Satz 3. Jede nichtleere, nach oben (bzw. unten) beschränkte Teilmenge $A \subset \mathbb{R}$ besitzt ein Supremum (bzw. Infimum).

Beweis. Sei $A \subset \mathbb{R}$ nichtleer und nach oben beschränkt. Dann gibt es ein Element $x_0 \in A$ und eine obere Schranke K_0 von A . Wir konstruieren jetzt durch Induktion nach n eine Folge von Intervallen

$$[x_0, K_0] \supset [x_1, K_1] \supset \dots \supset [x_n, K_n] \supset [x_{n+1}, K_{n+1}] \supset \dots$$

so dass für alle n gilt:

- (1) $x_n \in A$,
- (2) K_n ist obere Schranke von A ,
- (3) $K_n - x_n \leq 2^{-n}(K_0 - x_0)$.

Für $n = 0$ haben wir die Wahl von x_0 und K_0 schon durchgeführt.

Induktions-Schritt $n \rightarrow n + 1$. Sei $M := (K_n + x_n)/2$ die Mitte des Intervalls $[x_n, K_n]$. Es können zwei Fälle auftreten:

1. Fall: $A \cap]M, K_n] = \emptyset$. Dann ist M eine obere Schranke von A und wir definieren $x_{n+1} := x_n$ und $K_{n+1} := M$.

2. Fall: $A \cap]M, K_n] \neq \emptyset$. Dann gibt es einen Punkt $x_{n+1} \in A$ mit $x_{n+1} > M$. In diesem Fall setzen wir $K_{n+1} := K_n$.

In jedem der beiden Fälle gelten für x_{n+1} und K_{n+1} wieder die Eigenschaften (1) bis (3).

Nun ist $(K_n)_{n \in \mathbb{N}}$ eine monoton fallende und nach unten beschränkte Folge, konvergiert also nach §5, Satz 7 gegen eine Zahl K . Wir zeigen, dass K kleinste obere Schranke von A ist.

i) Für jedes $x \in A$ gilt $x \leq K_n$ für alle n . Daraus folgt $x \leq K$. Dies zeigt, dass K obere Schranke von A ist.

ii) Sei K' eine weitere obere Schranke von A . Angenommen, es würde gelten $K' < K$. Da $K - K' > 0$, gibt es ein n , so dass

$$K_n - x_n \leq 2^{-n}(K_0 - x_0) < K - K' \leq K_n - K'.$$

Dann folgt $K' < x_n$, was im Widerspruch dazu steht, dass K' obere Schranke von A ist. Also muss $K \leq K'$ sein, d.h. K ist kleinste obere Schranke von A . Wir haben somit die Existenz des Supremums von A bewiesen.

Die Existenz des Infimums zeigt man analog.

Beispiele

(9.6) Für das abgeschlossene Intervall $[a, b]$, $a \leq b$, gilt

$$\sup([a, b]) = b \quad \text{und} \quad \inf([a, b]) = a.$$

(9.7) Für das offene Intervall $]a, b[$, $a < b$, gilt ebenfalls

$$\sup(]a, b[) = b \quad \text{und} \quad \inf(]a, b[) = a.$$

Wir beweisen, dass b kleinste obere Schranke von $]a, b[$ ist. Zunächst ist klar, dass b obere Schranke ist. Um zu zeigen, dass b sogar kleinste obere Schranke ist, betrachten wir irgend eine obere Schranke K des Intervalls. Die Punkte $x_n := b - 2^{-n}(b - a)$ liegen für $n \geq 1$ alle im Intervall $]a, b[$, also ist $x_n \leq K$. Da $\lim_{n \rightarrow \infty} x_n = b$, folgt $b \leq K$. Also ist b kleinste obere Schranke.

(9.8) Für $A := \left\{ \frac{n}{n+1} : n \in \mathbb{N} \right\}$ gilt $\sup(A) = 1$.

(9.9) Für $A := \left\{ \frac{n^2}{2^n} : n \in \mathbb{N} \right\}$ gilt $\sup(A) = \frac{9}{8}$.

Beweis. $\frac{9}{8}$ ist obere Schranke von A , denn

$$\frac{n^2}{2^n} \leq 1 < \frac{9}{8} \quad \text{für alle } n \neq 3,$$

vgl. Aufgabe 3.1, und $\frac{3^2}{2^3} = \frac{9}{8}$. Außerdem ist $\frac{9}{8}$ kleinste obere Schranke, da $\frac{9}{8} \in A$.

Maximum, Minimum. Wie diese Beispiele zeigen, kann es sowohl vorkommen, dass $\sup(A)$ in A liegt, als auch, dass $\sup(A)$ nicht in A liegt. Falls $\sup(A) \in A$, nennt man $\sup(A)$ auch das *Maximum* von A . Ebenso heißt $\inf(A)$ das *Minimum* von A , falls $\inf(A) \in A$.

In jedem Fall existiert jedoch, wie aus dem Beweis von Satz 3 hervorgeht, eine Folge $x_n \in A$, $n \in \mathbb{N}$, mit $\lim_{n \rightarrow \infty} x_n = \sup(A)$ und eine Folge $y_n \in A$, $n \in \mathbb{N}$, mit $\lim_{n \rightarrow \infty} y_n = \inf(A)$, d.h. $\sup(A)$ und $\inf(A)$ sind Berührungspunkte von A .

Bezeichnung. Falls die Teilmenge $A \subset \mathbb{R}$ nicht nach oben (bzw. nicht nach unten) beschränkt ist, schreibt man

$$\sup(A) = +\infty \quad \text{bzw.} \quad \inf(A) = -\infty.$$

Limes superior, Limes inferior

Definition. Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge reeller Zahlen. Dann definiert man

$$\begin{aligned} \limsup_{n \rightarrow \infty} a_n &:= \lim_{n \rightarrow \infty} (\sup\{a_k : k \geq n\}), \\ \liminf_{n \rightarrow \infty} a_n &:= \lim_{n \rightarrow \infty} (\inf\{a_k : k \geq n\}). \end{aligned}$$

Eine andere Schreibweise ist $\overline{\lim}$ für $\limsup$ und $\underline{\lim}$ für $\liminf$.

Bemerkung. Die Folge $(\sup\{a_k : k \geq n\})_{n \in \mathbb{N}}$ ist monoton fallend (oder identisch $+\infty$) und die Folge $(\inf\{a_k : k \geq n\})_{n \in \mathbb{N}}$ ist monoton wachsend (oder identisch $-\infty$). Daher existieren $\limsup a_n$ und $\liminf a_n$ immer eigentlich oder uneigentlich, d.h. sie sind entweder reelle Zahlen oder es gilt $\limsup a_n = \pm\infty$ bzw. $\liminf a_n = \pm\infty$.

Beispiele

(9.10) Wir betrachten die Folge $a_n := (-1)^n(1 + \frac{1}{n})$, $n \geq 1$. Hier ist

$$\sup\{a_k : k \geq n\} = \begin{cases} 1 + \frac{1}{n}, & \text{falls } n \text{ gerade,} \\ 1 + \frac{1}{n+1}, & \text{falls } n \text{ ungerade.} \end{cases}$$

Also gilt $\limsup_{n \rightarrow \infty} a_n = 1$.

Entsprechend hat man

$$\inf\{a_k : k \geq n\} = \begin{cases} -(1 + \frac{1}{n}), & \text{falls } n \text{ ungerade,} \\ -(1 + \frac{1}{n+1}), & \text{falls } n \text{ gerade.} \end{cases}$$

Daraus folgt $\liminf_{n \rightarrow \infty} a_n = -1$.

(9.11) Für die Folge $a_n := n$, $n \in \mathbb{N}$, gilt

$$\begin{aligned} \sup\{a_k : k \geq n\} &= \infty, \\ \inf\{a_k : k \geq n\} &= n. \end{aligned}$$

Daraus folgt $\limsup_{n \rightarrow \infty} a_n = \infty$ und $\liminf_{n \rightarrow \infty} a_n = \infty$.

Satz 4 (Charakterisierung des Limes superior). *Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge reeller Zahlen und $a \in \mathbb{R}$. Genau dann gilt*

$$\limsup_{n \rightarrow \infty} a_n = a,$$

wenn für jedes $\varepsilon > 0$ die folgenden beiden Bedingungen erfüllt sind:

i) *Für fast alle Indizes $n \in \mathbb{N}$ (d.h. alle bis auf endlich viele) gilt*

$$a_n < a + \varepsilon.$$

ii) *Es gibt unendlich viele Indizes $m \in \mathbb{N}$ mit*

$$a_m > a - \varepsilon.$$

Beweis. Wir verwenden die Bezeichnungen

$$A_n := \{a_k : k \geq n\} \quad \text{und} \quad s_n := \sup A_n.$$

a) Sei zunächst vorausgesetzt, dass $\limsup a_n = a$, d.h. $\lim_{n \rightarrow \infty} s_n = a$, und sei $\varepsilon > 0$ vorgegeben. Da die Folge (s_n) monoton fällt, gilt $s_n \geq a$ für alle n . Daraus folgt Bedingung ii). Andererseits gibt es ein $N \in \mathbb{N}$, so dass $s_n < a + \varepsilon$ für alle $n \geq N$. Daraus folgt $a_n < a + \varepsilon$ für alle $n \geq N$.

b) Seien umgekehrt die Bedingungen i) und ii) erfüllt. Aus ii) folgt, dass $s_n > a - \varepsilon$ für alle n und alle $\varepsilon > 0$, also $\lim_{n \rightarrow \infty} s_n \geq a$. Wegen i) gibt es zu $\varepsilon > 0$ ein $N \in \mathbb{N}$, so dass $a_n < a + \varepsilon$ für alle $n \geq N$, woraus folgt $s_n \leq a + \varepsilon$ für $n \geq N$. Insgesamt folgt $\lim s_n = a$, q.e.d.

Bemerkung. Analog zu Satz 4 hat man folgende Charakterisierung des Limes inferior:

Es gilt $\liminf_{n \rightarrow \infty} a_n = a$ genau dann, wenn für jedes $\varepsilon > 0$ gilt:

- i) $a_n > a - \varepsilon$ für fast alle n , und
- ii) $a_n < a + \varepsilon$ für unendlich viele $n \in \mathbb{N}$.

AUFGABEN

9.1. Eine Zahl $x \in \mathbb{R}$ heißt *algebraisch*, wenn es eine natürliche Zahl $n \geq 1$ und rationale Zahlen $a_1, a_2, \dots, a_n \in \mathbb{Q}$ gibt, so dass

$$x^n + a_1 x^{n-1} + \dots + a_{n-1} x + a_n = 0.$$

Man beweise: Die Menge $A \subset \mathbb{R}$ aller algebraischen Zahlen ist abzählbar.

Hinweis. Man zeige dazu, dass die Menge aller Polynome mit rationalen Koeffizienten abzählbar ist und benutze (ohne Beweis), dass ein Polynom n -ten Grades höchstens n Nullstellen hat.

Bemerkung. Eine reelle Zahl $x \in \mathbb{R}$ heißt *transzendent*, wenn sie nicht algebraisch ist. Aus der Abzählbarkeit der algebraischen Zahlen folgt, dass es überabzählbar viele transzendente Zahlen gibt. Es ist jedoch i.Allg. schwer, von einer konkret gegebenen Zahl $x \in \mathbb{R}$ zu entscheiden, ob sie transzendent oder algebraisch ist. Bekannte transzendente Zahlen sind die Eulersche Zahl e (Beweis von Hermite 1873), sowie die Zahl π (Lindemann 1882). Aus der Transzendenz von π folgt die Unmöglichkeit der Quadratur des Kreises mit Zirkel und Lineal. Siehe dazu auch [T].

9.2. Man beweise:

- a) Die Menge aller *endlichen* Teilmengen von $\mathbb{N}$ ist abzählbar.
- b) Die Menge *aller* Teilmengen von $\mathbb{N}$ ist überabzählbar.

9.3. Man zeige, dass die Abbildung

$$\tau: \mathbb{N} \times \mathbb{N} \longrightarrow \mathbb{N}, \quad \tau(n, m) := \frac{1}{2}(n + m + 1)(n + m) + n,$$

bijektiv ist.

9.4. Man konstruiere bijektive Abbildungen

- i) $\phi_1 : \mathbb{R}^* \rightarrow \mathbb{R}$,
- ii) $\phi_2 : \mathbb{R} \setminus \mathbb{Z} \rightarrow \mathbb{R}$,
- iii) $\phi_3 : \mathbb{R}_+^* \rightarrow \mathbb{R}$,
- iv) $\phi_4 : \mathbb{R}_+ \rightarrow \mathbb{R}$.

Bemerkung. Eine Menge M heißt von der *Mächtigkeit des Kontinuums*, falls es eine bijektive Abbildung $\phi : M \rightarrow \mathbb{R}$ gibt. Die sog. *Kontinuums-Hypothese* von G. Cantor sagt, dass sich eine unendliche Teilmenge $M \subset \mathbb{R}$ entweder bijektiv auf $\mathbb{N}$ oder bijektiv auf $\mathbb{R}$ abbilden lässt. K. Gödel bewies 1938, dass die Kontinuums-Hypothese mit den üblichen Axiomen der Mengenlehre verträglich ist. P.J. Cohen bewies 1963, dass auch die Negation der Kontinuums-Hypothese (es gibt unendliche Teilmengen M von $\mathbb{R}$, die sich weder bijektiv auf $\mathbb{N}$ noch bijektiv auf $\mathbb{R}$ abbilden lassen), mit den üblichen Axiomen der Mengenlehre verträglich ist. Bei der Kontinuums-Hypothese handelt es sich also um ein unentscheidbares Problem.

9.5. Es sei $a \in \mathbb{R}_+^*$ und k eine natürliche Zahl ≥ 2 . Man zeige

$$\sup\{x \in \mathbb{Q} : x^k < a\} = \sqrt[k]{a}.$$

9.6. Sei $(a_n)_{n \in \mathbb{N}}$ eine beschränkte Folge reeller Zahlen. Man zeige: Genau dann gilt $\limsup_{n \rightarrow \infty} a_n = a$, wenn folgende beiden Bedingungen erfüllt sind:

- a) Es gibt eine konvergente Teilfolge $(a_{n_k})_{k \in \mathbb{N}}$ von (a_n) mit $\lim_{k \rightarrow \infty} a_{n_k} = a$.
- b) Für jede konvergente Teilfolge $(a_{m_k})_{k \in \mathbb{N}}$ von (a_n) gilt $\lim_{k \rightarrow \infty} a_{m_k} \leq a$.

9.7. Sei $(a_n)_{n \in \mathbb{N}}$ eine beschränkte Folge reeller Zahlen und H die Menge ihrer Häufungspunkte. Man zeige

$$\begin{aligned} \limsup_{n \rightarrow \infty} a_n &= \sup H, \\ \liminf_{n \rightarrow \infty} a_n &= \inf H. \end{aligned}$$

9.8. Man beweise: Eine Folge $(a_n)_{n \in \mathbb{N}}$ reeller Zahlen konvergiert genau dann gegen ein $a \in \mathbb{R}$, wenn

$$\limsup_{n \rightarrow \infty} a_n = \liminf_{n \rightarrow \infty} a_n = a.$$

9.9. Man untersuche, ob folgende Aussage richtig ist:

Eine Folge $(a_n)_{n \in \mathbb{N}}$ reeller Zahlen konvergiert genau dann gegen $a \in \mathbb{R}$, wenn für jede konvergente Teilfolge (a_{n_k}) von (a_n) gilt: $\lim_{k \rightarrow \infty} a_{n_k} = a$.

9.10. Man zeige, dass für jede Folge $(a_n)_{n \in \mathbb{N}}$ reeller Zahlen gilt:

$$\liminf_{n \rightarrow \infty} a_n = -\limsup_{n \rightarrow \infty} (-a_n).$$

9.11. Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge reeller Zahlen. Man zeige:

- Es gilt $\limsup_{n \rightarrow \infty} a_n = +\infty$ genau dann, wenn die Folge (a_n) nicht nach oben beschränkt ist.
- Es gilt $\limsup_{n \rightarrow \infty} a_n = -\infty$ genau dann, wenn die Folge (a_n) bestimmt gegen $-\infty$ divergiert.

9.12. Es seien $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ Folgen reeller Zahlen mit $\limsup a_n \neq -\infty$ und $\liminf b_n \neq -\infty$. Man zeige:

$$\limsup_{n \rightarrow \infty} a_n + \liminf_{n \rightarrow \infty} b_n \leq \limsup_{n \rightarrow \infty} (a_n + b_n) \leq \limsup_{n \rightarrow \infty} a_n + \limsup_{n \rightarrow \infty} b_n.$$

Dabei werde vereinbart $a + \infty = \infty + a = \infty$ für alle $a \in \mathbb{R} \cup \{\infty\}$.

9.13. Das *Dedekindsche Schnittaxiom* für einen angeordneten Körper K lautet wie folgt:

Seien $A, B \subset K$ nicht-leere Teilmengen mit $A \cup B = K$, so dass für alle $x \in A$ und $y \in B$ gilt $x < y$. Dann gibt es genau ein $s \in K$ mit

$$x \leq s \leq y \quad \text{für alle } x \in A \text{ und } y \in B.$$

Man beweise:

- Im Körper $\mathbb{R}$ gilt das Dedekindsche Schnittaxiom.
- In einem angeordneten Körper impliziert das Dedekindsche Schnittaxiom das Archimedische Axiom und das Vollständigkeits-Axiom.

§ 10 Funktionen. Stetigkeit

Wir kommen jetzt zu einem weiteren zentralen Begriff der Analysis, dem der stetigen Funktion. Wir zeigen, dass Summe, Produkt und Quotient (mit nichtverschwindendem Nenner) stetiger Funktionen sowie die Komposition stetiger Funktionen wieder stetig ist.

Definition. Sei D eine Teilmenge von $\mathbb{R}$. Unter einer reellwertigen (reellen) Funktion auf D versteht man eine Abbildung $f:D \rightarrow \mathbb{R}$. Die Menge D heißt Definitionsbereich von f . Der Graph von f ist die Menge

$$\Gamma_f := \{(x, y) \in D \times \mathbb{R} : y = f(x)\}.$$

Beispiele

(10.1) Konstante Funktionen. Sei $c \in \mathbb{R}$ vorgegeben.

$$\begin{aligned} f:\mathbb{R} &\rightarrow \mathbb{R}, \\ x &\mapsto f(x) = c. \end{aligned}$$

(10.2) Identische Abbildung

$$\begin{aligned} \text{id}_{\mathbb{R}}:\mathbb{R} &\rightarrow \mathbb{R}, \\ x &\mapsto x. \end{aligned}$$

(10.3) Absolutbetrag (Bild 10.1).

$$\begin{aligned} \text{abs}:\mathbb{R} &\rightarrow \mathbb{R}, \\ x &\mapsto |x|. \end{aligned}$$

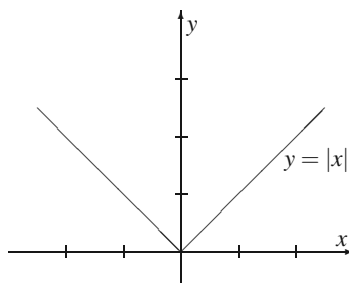


Bild 10.1 Absolutbetrag

(10.4) Die floor-Funktion (Bild 10.2).

$$\begin{aligned}\text{floor}: \mathbb{R} &\rightarrow \mathbb{R}, \\ x &\mapsto \lfloor x \rfloor.\end{aligned}$$

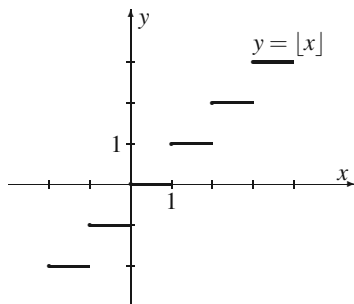


Bild 10.2 floor-Funktion

(10.5) Quadratwurzel (Bild 10.3).

$$\begin{aligned}\text{sqrt}: \mathbb{R}_+ &\rightarrow \mathbb{R}, \\ x &\mapsto \sqrt{x}.\end{aligned}$$

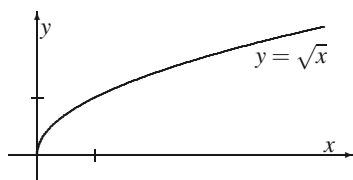


Bild 10.3 Quadratwurzel

(10.6) Exponentialfunktion (Bild 10.4).

$$\begin{aligned}\text{exp}: \mathbb{R} &\rightarrow \mathbb{R}, \\ x &\mapsto \exp(x).\end{aligned}$$

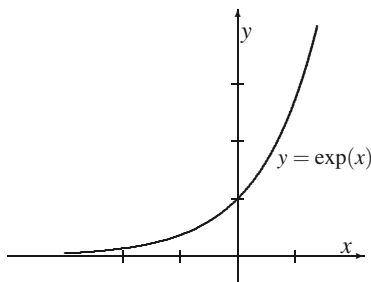


Bild 10.4 Exponentialfunktion

(10.7) Polynomfunktionen. Seien $a_0, a_1, \dots, a_n \in \mathbb{R}$.

$$\begin{aligned}p: \mathbb{R} &\rightarrow \mathbb{R}, \\ x &\mapsto p(x) := a_n x^n + \dots + a_1 x + a_0.\end{aligned}$$

(10.8) Rationale Funktionen. Seien

$$P(x) = a_n x^n + \dots + a_1 x + a_0,$$

$$Q(x) = b_m x^m + \dots + b_1 x + b_0$$

Polynome und $D := \{x \in \mathbb{R} : q(x) \neq 0\}$. Dann ist die rationale Funktion $R = \frac{P}{Q}$ definiert durch

$$R: D \rightarrow \mathbb{R}, \quad x \mapsto R(x) = \frac{P(x)}{Q(x)}.$$

Die Polynomfunktionen sind spezielle rationale Funktionen.

(10.9) Treppenfunktionen (Bild 10.5). Seien $a < b$ reelle Zahlen. Eine Funktion $\varphi: [a, b] \rightarrow \mathbb{R}$ heißt *Treppenfunktion*, wenn es eine Unterteilung

$$a = t_0 < t_1 < \dots < t_{n-1} < t_n = b$$

des Intervalls $[a, b]$ und Konstanten $c_1, c_2, \dots, c_n \in \mathbb{R}$ gibt, so dass

$$\varphi(x) = c_k \quad \text{für alle } x \in]t_{k-1}, t_k[, \quad (1 \leq k \leq n).$$

Die Funktionswerte $\varphi(t_k)$ in den Teilpunkten t_k sind beliebig.

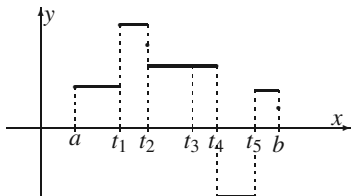


Bild 10.5 Treppenfunktion

(10.10) Es gibt auch Funktionen, deren Graph man nicht zeichnen kann, z.B. die von Dirichlet betrachtete Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ mit

$$f(x) = \begin{cases} 1, & \text{falls } x \text{ rational,} \\ 0, & \text{falls } x \text{ irrational.} \end{cases}$$

Definition (Rationale Operationen auf Funktionen). Seien $f, g: D \rightarrow \mathbb{R}$ Funktionen und $\lambda \in \mathbb{R}$. Dann sind die Funktionen

$$f + g: D \rightarrow \mathbb{R},$$

$$\lambda f: D \rightarrow \mathbb{R},$$

$$fg: D \rightarrow \mathbb{R}$$

definiert durch

$$(f + g)(x) := f(x) + g(x),$$

$$(\lambda f)(x) := \lambda f(x),$$

$$(fg)(x) := f(x)g(x).$$

Sei $D' := \{x \in D : g(x) \neq 0\}$. Dann ist die Funktion

$$\frac{f}{g}: D' \rightarrow \mathbb{R}$$

definiert durch

$$\left(\frac{f}{g}\right)(x) := \frac{f(x)}{g(x)}.$$

Bemerkung. Alle rationalen Funktionen entstehen aus $\text{id}_{\mathbb{R}}$ und der konstanten Funktion 1 durch wiederholte Anwendung dieser Operationen.

Die nächste Definition gibt ein weiteres Verfahren an, aus gegebenen Funktionen neue zu konstruieren.

Definition (Komposition von Funktionen). Seien $f: D \rightarrow \mathbb{R}$ und $g: E \rightarrow \mathbb{R}$ Funktionen mit $f(D) \subset E$. Dann ist die Funktion

$$g \circ f: D \rightarrow \mathbb{R}$$

definiert durch $(g \circ f)(x) := g(f(x))$ für $x \in D$.

(10.11) Beispiel. Sei $q: \mathbb{R} \rightarrow \mathbb{R}$ die durch $q(x) = x^2$ definierte Funktion. Dann läßt sich die Funktion $\text{abs}: \mathbb{R} \rightarrow \mathbb{R}$ schreiben als

$$\text{abs} = \text{sqrt} \circ q,$$

denn es gilt für alle $x \in \mathbb{R}$

$$(\text{sqrt} \circ q)(x) = \text{sqrt}(q(x)) = \sqrt{x^2} = |x| = \text{abs}(x).$$

Grenzwerte bei Funktionen

Wir verbinden jetzt den Grenzwertbegriff und den Funktionsbegriff.

Definition. Sei $f: D \rightarrow \mathbb{R}$ eine reelle Funktion auf $D \subset \mathbb{R}$ und $a \in \mathbb{R}$ ein Berührungspunkt von D . Man definiert

$$\lim_{x \rightarrow a} f(x) = c,$$

falls für *jede* Folge $(x_n)_{n \in \mathbb{N}}$, $x_n \in D$, mit $\lim_{n \rightarrow \infty} x_n = a$ gilt:

$$\lim_{n \rightarrow \infty} f(x_n) = c.$$

Statt $\lim_{x \rightarrow a} f(x)$ schreibt man zur Verdeutlichung auch $\lim_{\substack{x \rightarrow a \\ x \in D}} f(x)$.

Bemerkung. Da a Berührungspunkt von D ist (vgl. die Definition in § 9, Seite 94), gibt es mindestens eine Folge (x_n) , $x_n \in D$, mit $\lim_{n \rightarrow \infty} x_n = a$.

Falls $a \in D$, hat man in jedem Fall die konstante Folge $x_n = a$ für alle n , so dass dann der Limes $\lim_{x \rightarrow a} f(x)$ im Falle der Existenz notwendig gleich $f(a)$ ist.

Weitere Bezeichnungen

a) $\lim_{x \searrow a} f(x) = c$ bedeutet:

a ist Berührungspunkt von $D \cap]a, \infty[$ und für jede Folge (x_n) mit $x_n \in D$, $x_n > a$ und $\lim_{n \rightarrow \infty} x_n = a$ gilt

$$\lim_{n \rightarrow \infty} f(x_n) = c.$$

b) $\lim_{x \nearrow a} f(x) = c$ bedeutet:

a ist Berührungspunkt von $D \cap]-\infty, a[$ und für jede Folge (x_n) mit $x_n \in D$, $x_n < a$ und $\lim_{n \rightarrow \infty} x_n = a$ gilt

$$\lim_{n \rightarrow \infty} f(x_n) = c.$$

c) $\lim_{x \rightarrow \infty} f(x) = c$ bedeutet:

Der Definitionsbereich D ist nach oben unbeschränkt und für jede Folge (x_n) mit $x_n \in D$ und $\lim_{n \rightarrow \infty} x_n = \infty$ gilt

$$\lim_{n \rightarrow \infty} f(x_n) = c.$$

Analog ist $\lim_{x \rightarrow -\infty} f(x)$ definiert.

Beispiele

(10.12) $\lim_{x \rightarrow 0} \exp(x) = 1$.

Beweis. Die Restgliedabschätzung aus § 8, Satz 2, liefert für $N = 0$:

$$|\exp(x) - 1| \leq 2|x| \quad \text{für } |x| \leq 1.$$

Sei (x_n) eine beliebige Folge mit $\lim x_n = 0$. Dann gilt $|x_n| < 1$ für alle $n \geq n_0$, also

$$|\exp(x_n) - 1| \leq 2|x_n| \quad \text{für } n \geq n_0.$$

Daraus folgt $\lim_{n \rightarrow \infty} |\exp(x_n) - 1| = 0$, also

$$\lim_{n \rightarrow \infty} \exp(x_n) = 1, \quad \text{q.e.d.}$$

(10.13) Es gilt $\lim_{x \searrow 1} \text{floor}(x) = 1$ und $\lim_{x \nearrow 1} \text{floor}(x) = 0$;

Also existiert $\lim_{x \rightarrow 1} \text{floor}(x)$ nicht.

(10.14) Es sei $P: \mathbb{R} \rightarrow \mathbb{R}$ ein Polynom der Gestalt

$$P(x) = x^k + a_1 x^{k-1} + \dots + a_{k-1} x + a_k, \quad (k \geq 1).$$

Dann gilt

$$\begin{aligned} \lim_{x \rightarrow \infty} P(x) &= \infty, \\ \lim_{x \rightarrow -\infty} P(x) &= \begin{cases} +\infty, & \text{falls } k \text{ gerade,} \\ -\infty, & \text{falls } k \text{ ungerade.} \end{cases} \end{aligned}$$

Beweis. Für $x \neq 0$ gilt $P(x) = x^k g(x)$, wobei

$$g(x) = 1 + \frac{a_1}{x} + \frac{a_2}{x^2} + \dots + \frac{a_k}{x^k}.$$

Für alle $x \in \mathbb{R}$ mit

$$x \geq c := \max(1, 2k|a_1|, 2k|a_2|, \dots, 2k|a_k|)$$

gilt $g(x) \geq \frac{1}{2}$, also $P(x) \geq \frac{1}{2}x^k \geq \frac{x}{2}$. Sei nun (x_n) eine beliebige Folge reeller Zahlen mit $\lim x_n = \infty$. Dann gilt $x_n \geq c$ für alle $n \geq n_0$, also $P(x_n) \geq \frac{1}{2}x_n$ für $n \geq n_0$. Daraus folgt

$$\lim_{n \rightarrow \infty} P(x_n) = \infty.$$

Die Behauptung über den Limes für $x \rightarrow -\infty$ folgt aus der Tatsache, dass

$$P(-x) = (-1)^k Q(x)$$

mit

$$Q(x) = x^k - a_1 x^{k-1} + \dots + (-1)^{k-1} a_{k-1} x + (-1)^k a_k.$$

Stetige Funktionen

Definition. Sei $f: D \rightarrow \mathbb{R}$ eine Funktion und $a \in D$. Die Funktion f heißt *stetig* im Punkt a , falls

$$\lim_{x \rightarrow a} f(x) = f(a).$$

f heißt stetig in D , falls f in jedem Punkt von D stetig ist.

Beispiele

(10.15) Die konstanten Funktionen und $\text{id}_{\mathbb{R}}$ sind überall stetig.

(10.16) Die Exponentialfunktion $\exp: \mathbb{R} \rightarrow \mathbb{R}$ ist in jedem Punkt stetig.

Beweis. Sei $a \in \mathbb{R}$. Wir haben zu zeigen, dass

$$\lim_{x \rightarrow a} \exp(x) = \exp(a).$$

Sei (x_n) eine beliebige Folge mit $\lim x_n = a$. Dann gilt $\lim(x_n - a) = 0$, also nach Beispiel (10.12)

$$\lim_{n \rightarrow \infty} \exp(x_n - a) = 1.$$

Daraus folgt mithilfe der Funktionalgleichung

$$\begin{aligned} \lim_{n \rightarrow \infty} \exp(x_n) &= \lim_{n \rightarrow \infty} (\exp(a) \exp(x_n - a)) \\ &= \exp(a) \lim_{n \rightarrow \infty} \exp(x_n - a) = \exp(a), \quad \text{q.e.d.} \end{aligned}$$

(10.17) Die in Beispiel (10.10) definierte Dirichletsche Funktion ist in keinem Punkt $x \in \mathbb{R}$ stetig.

Satz 1. Seien $f, g: D \rightarrow \mathbb{R}$ Funktionen, die in $a \in D$ stetig sind und sei $\lambda \in \mathbb{R}$. Dann sind auch die Funktionen

$$f + g: D \rightarrow \mathbb{R},$$

$$\lambda f: D \rightarrow \mathbb{R},$$

$$fg: D \rightarrow \mathbb{R}$$

im Punkte a stetig. Ist $g(a) \neq 0$, so ist auch die Funktion

$$\frac{f}{g}: D' \rightarrow \mathbb{R}$$

in a stetig. Dabei ist $D' = \{x \in D : g(x) \neq 0\}$.

Beweis. Sei (x_n) eine Folge $x_n \in D$ (bzw. $x_n \in D'$) und $\lim x_n = a$. Es ist zu zeigen:

$$\lim_{n \rightarrow \infty} (f+g)(x_n) = (f+g)(a),$$

$$\lim_{n \rightarrow \infty} (\lambda f)(x_n) = (\lambda f)(a),$$

$$\lim_{n \rightarrow \infty} (fg)(x_n) = (fg)(a), \quad \lim_{n \rightarrow \infty} \left(\frac{f}{g} \right)(x_n) = \left(\frac{f}{g} \right)(a).$$

Nach Voraussetzung ist $\lim_{n \rightarrow \infty} f(x_n) = f(a)$ und $\lim_{n \rightarrow \infty} g(x_n) = g(a)$. Die Behauptung folgt deshalb aus den in §4, Satz 3 bis 4, aufgestellten Rechenregeln für Zahlenfolgen.

Corollar. *Alle rationalen Funktionen sind stetig in ihrem Definitionsbereich.*

Dies folgt durch wiederholte Anwendung von Satz 1 auf Beispiel (10.15).

Bemerkung. Die Stetigkeit ist eine lokale Eigenschaft in folgendem Sinn: Seien $f, g: D \rightarrow \mathbb{R}$ zwei Funktionen, die in einer Umgebung eines Punktes $a \in D$ übereinstimmen, d.h. es gebe ein $\varepsilon > 0$, so dass $f(x) = g(x)$ für alle $x \in D$ mit $|x - a| < \varepsilon$. Dann ist f genau dann in a stetig, wenn g in a stetig ist. Dies folgt unmittelbar aus der Definition.

(10.18) Beispiel. Die Funktion $\text{abs}: \mathbb{R} \rightarrow \mathbb{R}$ ist stetig.

Beweis. Sei a ein beliebiger Punkt aus $\mathbb{R}$.

1. Fall: $a > 0$. In der Umgebung $]0, 2a[$ von a gilt $\text{abs}(x) = x = \text{id}_{\mathbb{R}}(x)$. Da $\text{id}_{\mathbb{R}}$ stetig ist, ist auch abs in a stetig.

2. Fall: $a < 0$. In der Umgebung $]2a, 0[$ von a gilt $\text{abs}(x) = -x = -\text{id}_{\mathbb{R}}(x)$. Nach Satz 1 ist $-\text{id}_{\mathbb{R}}$ stetig, also ist auch abs im Punkt a stetig.

3. Fall: $a = 0$. Sei (x_n) eine Folge mit $\lim x_n = 0$. Dann ist

$$\lim(\text{abs}(x_n)) = \lim|x_n| = 0 = \text{abs}(0),$$

also abs in 0 stetig.

Satz 2 (Komposition stetiger Funktionen). *Seien $f: D \rightarrow \mathbb{R}$ und $g: E \rightarrow \mathbb{R}$ Funktionen mit $f(D) \subset E$. Die Funktion f sei in $a \in D$ und g in $b := f(a) \in E$ stetig. Dann ist die Funktion*

$$g \circ f: D \rightarrow \mathbb{R}$$

in a stetig.

Beweis. Sei (x_n) eine Folge mit $x_n \in D$ und $\lim x_n = a$. Wegen der Stetigkeit von f in a gilt $\lim_{n \rightarrow \infty} f(x_n) = f(a)$. Nach Voraussetzung ist $y_n := f(x_n) \in E$ und $\lim y_n = b$. Da g in b stetig ist, gilt $\lim_{n \rightarrow \infty} g(y_n) = g(b)$. Deshalb folgt

$$\begin{aligned} \lim_{n \rightarrow \infty} (g \circ f)(x_n) &= \lim_{n \rightarrow \infty} g(f(x_n)) = \lim_{n \rightarrow \infty} g(y_n) = g(b) \\ &= g(f(a)) = (g \circ f)(a), \quad \text{q.e.d.} \end{aligned}$$

Beispiele

(10.19) Sei $f: D \rightarrow \mathbb{R}$ stetig. Dann ist auch die Funktion

$$|f|: D \rightarrow \mathbb{R}, \quad x \mapsto |f(x)|$$

stetig. Denn es gilt $|f| = \text{abs} \circ f$.

(10.20) Wir gehen aus von den stetigen Funktionen

$$\exp: \mathbb{R} \rightarrow \mathbb{R},$$

$$q: \mathbb{R} \rightarrow \mathbb{R}, \quad x \mapsto x^2.$$

Nach Satz 2 sind auch die Zusammensetzungen $f := \exp \circ q$ und $\varphi := q \circ \exp$, d.h. die Funktionen

$$f: \mathbb{R} \rightarrow \mathbb{R}, \quad x \mapsto \exp(x^2)$$

und

$$\varphi: \mathbb{R} \rightarrow \mathbb{R}, \quad x \mapsto (\exp(x))^2$$

stetig.

AUFGABEN

10.1. Die Funktionen

$$\cosh: \mathbb{R} \rightarrow \mathbb{R} \quad (\text{Cosinus hyperbolicus}),$$

$$\sinh: \mathbb{R} \rightarrow \mathbb{R} \quad (\text{Sinus hyperbolicus})$$

sind definiert durch

$$\cosh(x) := \frac{1}{2}(\exp(x) + \exp(-x)),$$

$$\sinh(x) := \frac{1}{2}(\exp(x) - \exp(-x)).$$

Man zeige, dass diese Funktionen stetig sind und beweise für $x, y \in \mathbb{R}$ die Formeln

$$\cosh(x+y) = \cosh(x) \cosh(y) + \sinh(x) \sinh(y),$$

$$\sinh(x+y) = \cosh(x) \sinh(y) + \sinh(x) \cosh(y),$$

$$\cosh^2(x) - \sinh^2(x) = 1.$$

10.2. Die Funktionen $g_n: \mathbb{R} \rightarrow \mathbb{R}$, $n \in \mathbb{N}$, seien definiert durch

$$g_n(x) := \frac{nx}{1+|nx|}.$$

Man zeige, dass alle Funktionen g_n stetig sind. Für welche $x \in \mathbb{R}$ ist die Funktion

$$g(x) := \lim_{n \rightarrow \infty} g_n(x)$$

definiert bzw. stetig?

10.3. Die Funktion $\text{zack}: \mathbb{R} \rightarrow \mathbb{R}$ sei definiert durch

$$\text{zack}(x) := \text{abs}\left(\left\lfloor x + \frac{1}{2} \right\rfloor - x\right).$$

Man zeichne den Graphen der Funktion zack und zeige:

- Für $|x| \leq \frac{1}{2}$ gilt $\text{zack}(x) = \text{abs}(x)$.
- Für alle $x \in \mathbb{R}$ und $n \in \mathbb{Z}$ gilt $\text{zack}(x+n) = \text{zack}(x)$.
- Die Funktion zack ist stetig.

10.4. Für $p, q \in \mathbb{Z}$ und $x \in \mathbb{R}^*$ sei definiert

$$f_{pq}(x) := |x|^p \text{zack}(x^q).$$

Für welche p und q kann man $f_{pq}(0)$ so definieren, dass eine überall stetige Funktion $f_{pq}: \mathbb{R} \rightarrow \mathbb{R}$ entsteht?

10.5. Die Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ sei definiert durch

$$f(x) := \begin{cases} 1/q, & \text{falls } x = \pm p/q \text{ mit } p, q \in \mathbb{N} \text{ teilerfremd,} \\ 0, & \text{falls } x \text{ irrational.} \end{cases}$$

Man zeige, dass f in jedem Punkt $a \in \mathbb{R} \setminus \mathbb{Q}$ stetig ist.

10.6. Die Funktion $f: \mathbb{Q} \rightarrow \mathbb{R}$ werde definiert durch

$$f(x) := \begin{cases} 0, & \text{falls } x < \sqrt{2}, \\ 1, & \text{falls } x > \sqrt{2}. \end{cases}$$

Man zeige, dass f auf $\mathbb{Q}$ stetig ist.

§ 11 Sätze über stetige Funktionen

In diesem Paragraphen beweisen wir die wichtigsten allgemeinen Sätze über stetige Funktionen in abgeschlossenen und beschränkten Intervallen, nämlich den Zwischenwertsatz, den Satz über die Annahme von Maximum und Minimum und die gleichmäßige Stetigkeit.

Satz 1 (Zwischenwertsatz). *Sei $f: [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion mit $f(a) < 0$ und $f(b) > 0$ (bzw. $f(a) > 0$ und $f(b) < 0$). Dann existiert ein $p \in [a, b]$ mit $f(p) = 0$.*

Bemerkung. Die Aussage des Satzes ist anschaulich klar, vgl. Bild 11.1. Sie bedarf aber natürlich dennoch eines Beweises, da eine Zeichnung keinerlei Beweiskraft hat. Die Aussage wird falsch, wenn man nur innerhalb der rationalen Zahlen arbeitet. Sei etwa $D := \{x \in \mathbb{Q} : 1 \leq x \leq 2\}$ und $f: D \rightarrow \mathbb{R}$ die stetige Funktion $x \mapsto f(x) = x^2 - 2$. Dann ist $f(1) = -1 < 0$ und $f(2) = 2 > 0$, aber es gibt kein $p \in D$ mit $f(p) = 0$, da die Zahl 2 keine rationale Quadratwurzel hat.

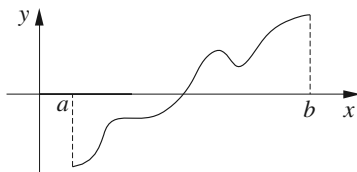


Bild 11.1

Beweis. Wir benutzen die Intervall-Halbierungsmethode. Sei $f(a) < 0$ und $f(b) > 0$. Wir definieren induktiv eine Folge $[a_n, b_n] \subset [a, b]$, $n \in \mathbb{N}$, von Intervallen mit folgenden Eigenschaften:

- (1) $[a_n, b_n] \subset [a_{n-1}, b_{n-1}]$ für $n \geq 1$,
- (2) $b_n - a_n = 2^{-n}(b - a)$,
- (3) $f(a_n) \leq 0$, $f(b_n) \geq 0$.

Induktionsanfang. Wir setzen $[a_0, b_0] := [a, b]$.

Induktionsschritt. Sei das Intervall $[a_n, b_n]$ bereits definiert und sei $m := (a_n + b_n)/2$ die Mitte des Intervalls. Nun können zwei Fälle auftreten:

1. Fall: $f(m) \geq 0$. Dann sei $[a_{n+1}, b_{n+1}] := [a_n, m]$.
2. Fall: $f(m) < 0$. Dann sei $[a_{n+1}, b_{n+1}] := [m, b_n]$.

Offenbar sind wieder die Eigenschaften (1)–(3) für $n+1$ erfüllt. Es folgt, dass die Folge (a_n) monoton wachsend und beschränkt und die Folge (b_n) monoton fallend und beschränkt ist. Also konvergieren beide Folgen (§5, Satz 7) und wegen (2) gilt

$$\lim_{n \rightarrow \infty} a_n = \lim_{n \rightarrow \infty} b_n =: p.$$

Aufgrund der Stetigkeit von f ist $\lim f(a_n) = \lim f(b_n) = f(p)$. Aus (3) folgt nach §4, Corollar zu Satz 5, dass

$$f(p) = \lim f(a_n) \leq 0 \quad \text{und} \quad f(p) = \lim f(b_n) \geq 0.$$

Daher gilt $f(p) = 0$, q.e.d.

(11.1) Beispiel. Jedes Polynom ungeraden Grades $f: \mathbb{R} \rightarrow \mathbb{R}$,

$$f(x) = x^n + c_1 x^{n-1} + \dots + c_n,$$

besitzt mindestens eine reelle Nullstelle.

Denn nach (10.14) gilt

$$\lim_{x \rightarrow -\infty} f(x) = -\infty \quad \text{und} \quad \lim_{x \rightarrow \infty} f(x) = \infty,$$

man kann also Stellen $a < b$ finden mit $f(a) < 0$ und $f(b) > 0$. Deshalb gibt es ein $p \in [a, b]$ mit $f(p) = 0$.

Bemerkung. Ein Polynom geraden Grades braucht keine reelle Nullstelle zu besitzen, wie das Beispiel $f(x) = x^{2k} + 1$ zeigt.

Das Intervallhalbierungs-Verfahren ist konstruktiv und kann auch zur praktischen Nullstellen-Berechnung verwendet werden. Zur Illustration schreiben wir eine kleine ARIBAS-Funktion `findzero`, die als Argumente eine Funktion `f` und zwei Stellen `a`, `b`, an denen die Funktion verschiedenes Vorzeichen hat, erwartet, sowie eine positive Fehlerschranke `eps`.

```
function findzero(f: function; a,b,eps: real): real;
var
  x1,x2,y1,y2,m: real;
begin
  y1 := f(a); y2 := f(b);
  if (y1 > 0 and y2 > 0) or (y1 < 0 and y2 < 0) then
    writeln("bad interval [a,b]");
    halt();
  elsif y1 < 0 then
    x1 := a; x2 := b;
  else
    x1 := b; x2 := a;
  end;
  while abs(x2-x1) > eps do
    m := (x1 + x2)/2;
    if f(m) >= 0 then
      x2 := m;
    else
      x1 := m;
    end;
  end;
  return (x1 + x2)/2;
end.
```

Die Funktion prüft zuerst die Vorzeichen von $f(a)$ und $f(b)$, und steigt mit Fehlermeldung aus, falls diese gleich sind. Je nachdem $f(a)$ negativ oder nicht-negativ ist, wird a der Variablen x_1 und b der Variablen x_2 zugeordnet, oder umgekehrt. Dann beginnt das Intervall-Halbierungsverfahren, bis die Länge des Intervalls kleiner-gleich der Fehlerschranke eps wird. Die Mitte des letzten Intervalls wird ausgegeben.

Wir testen `findzero` für die Funktion $f(x) := x^5 - x - 1$; man sieht unmittelbar, dass $f(0) < 0$ und $f(2) > 0$. Wir schreiben für f die ARIBAS-Funktion `testfun`.

```
function testfun(x: real): real
begin
  return x**5 - x - 1;
end.
```

Als Fehlerschranke wählen wir 10^{-7} .

```

==> eps := 10**-7.
-: 1.00000000E-7

==> x0 := findzero(testfun,0,2,eps) .
-: 1.16730395

```

Um zu verifizieren, dass damit eine Nullstelle mit der gewünschten Genauigkeit gefunden wurde, berechnen wir f an den Stellen $x_0 - \varepsilon/2$ und $x_0 + \varepsilon/2$.

```

==> testfun(x0 - eps/2) .
-: -6.50528818E-7

==> testfun(x0 + eps/2) .
-: 1.74622983E-7

```

Also haben wir tatsächlich eine Nullstelle von f bis auf einen Fehler $\pm 0.5 \cdot 10^{-7}$ gefunden. In diesem Zusammenhang sei noch auf ein Problem beim numerischen Rechnen hingewiesen: Ist $f(x)$ sehr nahe bei 0, so ist es wegen der Rechenungenauigkeit manchmal unmöglich, numerisch zu entscheiden, ob $f(x)$ größer, kleiner, oder gleich 0 ist. In unserem Beispiel tritt dieses Problem nicht auf.

Corollar 1. Sei $f: [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion und c eine reelle Zahl zwischen $f(a)$ und $f(b)$. Dann existiert ein $p \in [a, b]$ mit $f(p) = c$.

Beweis. Sei etwa $f(a) < c < f(b)$. Die Funktion $g: [a, b] \rightarrow \mathbb{R}$ sei definiert durch $g(x) := f(x) - c$. Dann ist g stetig und $g(a) < 0 < g(b)$. Nach Satz 1 existiert daher ein $p \in [a, b]$ mit $g(p) = 0$, woraus folgt $f(p) = c$, q.e.d.

Corollar 2. Sei $I \subset \mathbb{R}$ ein (eigentliches oder uneigentliches) Intervall und $f: I \rightarrow \mathbb{R}$ eine stetige Funktion. Dann ist auch $f(I) \subset \mathbb{R}$ ein Intervall.

Beweis. Wir setzen

$$B := \sup f(I) \in \mathbb{R} \cup \{+\infty\}, \quad A := \inf f(I) \in \mathbb{R} \cup \{-\infty\}$$

und zeigen zunächst, dass $]A, B[\subset f(I)$. Denn sei y irgend eine Zahl mit $A < y < B$. Nach Definition von A und B gibt es dann $a, b \in I$ mit $f(a) < y < f(b)$. Nach Corollar 1 existiert ein $x \in I$ mit $f(x) = y$; also ist $y \in f(I)$. Damit ist $]A, B[\subset f(I)$ bewiesen. Es folgt, dass $f(I)$ gleich einem der folgenden vier Intervalle ist: $]A, B[$, $]A, B]$, $[A, B[$ oder $[A, B]$.

Definition. Eine Funktion $f: D \rightarrow \mathbb{R}$ heißt *beschränkt*, wenn die Menge $f(D)$ beschränkt ist, d.h. wenn ein $M \in \mathbb{R}_+$ existiert, so dass

$$|f(x)| \leq M \quad \text{für alle } x \in D.$$

Definition. Unter einem *kompakten Intervall* versteht man ein abgeschlossenes und beschränktes Intervall $[a, b] \subset \mathbb{R}$.

Satz 2. Jede in einem kompakten Intervall stetige Funktion $f: [a, b] \rightarrow \mathbb{R}$ ist beschränkt und nimmt ihr Maximum und Minimum an, d.h. es existiert ein Punkt $p \in [a, b]$, so dass

$$f(p) = \sup\{f(x) : x \in [a, b]\}$$

und ein Punkt $q \in [a, b]$, so dass

$$f(q) = \inf\{f(x) : x \in [a, b]\}.$$

Bemerkung. Satz 2 gilt nicht in offenen, halboffenen oder uneigentlichen Intervallen. Z.B. ist die Funktion $f:]0, 1] \rightarrow \mathbb{R}$, $f(x) := 1/x$, in $]0, 1]$ stetig, aber nicht beschränkt. Die Funktion $g:]0, 1[\rightarrow \mathbb{R}$, $g(x) := x$, ist stetig und beschränkt, nimmt aber weder ihr Infimum 0 noch ihr Supremum 1 an.

Beweis. Wir geben nur den Beweis für das Maximum. Der Übergang von f zu $-f$ liefert dann die Behauptung für das Minimum. Sei

$$A := \sup\{f(x) : x \in [a, b]\} \in \mathbb{R} \cup \{\infty\}.$$

(Es gilt $A = \infty$, falls f nicht nach oben beschränkt ist.) Dann existiert eine Folge $x_n \in [a, b]$, $n \in \mathbb{N}$, so dass

$$\lim_{n \rightarrow \infty} f(x_n) = A.$$

Da die Folge (x_n) beschränkt ist, besitzt sie nach dem Satz von Bolzano-Weierstraß eine konvergente Teilfolge $(x_{n_k})_{k \in \mathbb{N}}$ mit

$$\lim_{k \rightarrow \infty} x_{n_k} =: p \in [a, b].$$

Aus der Stetigkeit von f folgt

$$f(p) = \lim_{k \rightarrow \infty} f(x_{n_k}) = A,$$

insbesondere $A \in \mathbb{R}$, also ist f nach oben beschränkt und nimmt in p ihr Maximum an.

Der folgende Satz gibt eine Umformulierung der Definition der Stetigkeit.

Satz 3 (ϵ - δ -Definition der Stetigkeit). *Sei $D \subset \mathbb{R}$ und $f: D \rightarrow \mathbb{R}$ eine Funktion. f ist genau dann im Punkt $p \in D$ stetig, wenn gilt:*

Zu jedem $\epsilon > 0$ existiert ein $\delta > 0$, so dass
 $|f(x) - f(p)| < \epsilon$ *für alle $x \in D$ mit $|x - p| < \delta$.*

Man kann dies in Worten auch so ausdrücken: f ist genau dann in p stetig, wenn gilt: Der Funktionswert $f(x)$ weicht beliebig wenig von $f(p)$ ab, falls nur x hinreichend nahe bei p liegt.

Beweis. 1) Es gebe zu jedem $\epsilon > 0$ ein $\delta > 0$, so dass $|f(x) - f(p)| < \epsilon$ für alle $x \in D$ mit $|x - p| < \delta$.

Es ist zu zeigen, dass für jede Folge (x_n) mit $x_n \in D$ und $\lim x_n = p$ gilt $\lim f(x_n) = f(p)$.

Sei $\epsilon > 0$ vorgegeben und sei $\delta > 0$ gemäß Voraussetzung. Wegen $\lim x_n = p$ existiert ein $N \in \mathbb{N}$, so dass $|x_n - p| < \delta$ für alle $n \geq N$.

Nach Voraussetzung ist daher $|f(x_n) - f(p)| < \epsilon$ für alle $n \geq N$. Also gilt $\lim_{n \rightarrow \infty} f(x_n) = f(p)$.

2) Für jede Folge $x_n \in D$ mit $\lim x_n = p$ gelte $\lim_{n \rightarrow \infty} f(x_n) = f(p)$. Es ist zu zeigen: Zu jedem $\epsilon > 0$ existiert ein $\delta > 0$, so dass

$$|f(x) - f(p)| < \epsilon \quad \text{für alle } x \in D \text{ mit } |x - p| < \delta.$$

Angenommen, dies sei nicht der Fall. Dann gibt es ein $\epsilon > 0$, so dass kein $\delta > 0$ existiert mit $|f(x) - f(p)| < \epsilon$ für alle $x \in D$ mit $|x - p| < \delta$. Es existiert also zu jedem $\delta > 0$ wenigstens ein $x \in D$ mit $|x - p| < \delta$, aber $|f(x) - f(p)| \geq \epsilon$. Insbesondere gibt es dann für jede natürliche Zahl $n \geq 1$ ein $x_n \in D$ mit

$$|x_n - p| < \frac{1}{n} \quad \text{und} \quad |f(x_n) - f(p)| \geq \epsilon.$$

Folglich ist $\lim x_n = p$ und daher nach Voraussetzung $\lim f(x_n) = f(p)$. Dies steht aber im Widerspruch zu $|f(x_n) - f(p)| \geq \epsilon$ für alle $n \geq 1$.

Corollar. *Sei $f: D \rightarrow \mathbb{R}$ stetig im Punkt $p \in D$ und $f(p) \neq 0$. Dann ist $f(x) \neq 0$ für alle x in einer Umgebung von p , d.h. es existiert ein $\delta > 0$, so dass*

$$f(x) \neq 0 \quad \text{für alle } x \in D \text{ mit } |x - p| < \delta.$$

Beweis. Zu $\epsilon := |f(p)| > 0$ existiert nach Satz 3 ein $\delta > 0$, so dass

$$|f(x) - f(p)| < \epsilon \quad \text{für alle } x \in D \text{ mit } |x - p| < \delta.$$

Daraus folgt $|f(x)| \geq |f(p)| - |f(x) - f(p)| > 0$ für alle $x \in D$ mit $|x - p| < \delta$, q.e.d.

Gleichmäßige Stetigkeit

Wir kommen jetzt zu einem wichtigen Begriff, der eine Verschärfung des Begriffs der Stetigkeit darstellt.

Definition. Eine Funktion $f: D \rightarrow \mathbb{R}$ heißt in D *gleichmäßig stetig*, wenn gilt:

Zu jedem $\varepsilon > 0$ existiert ein $\delta > 0$, so dass

$$|f(x) - f(x')| < \varepsilon \quad \text{für alle } x, x' \in D \text{ mit } |x - x'| < \delta.$$

Bemerkung. Vergleicht man dies mit der ε - δ -Definition der Stetigkeit aus Satz 3, so sieht man, dass eine gleichmäßig stetige Funktion $f: D \rightarrow \mathbb{R}$ in jedem Punkt $p \in D$ stetig ist. Der Unterschied beider Definitionen ist, dass bei gleichmäßiger Stetigkeit das δ nur von ε , aber nicht vom Punkt p abhängen darf. Für stetige Funktionen auf *kompakten* Intervallen läuft dies aber auf dasselbe hinaus, wie der folgende Satz zeigt.

Satz 4. Jede auf einem kompakten Intervall stetige Funktion $f: [a, b] \rightarrow \mathbb{R}$ ist dort gleichmäßig stetig.

Beweis. Angenommen, f sei nicht gleichmäßig stetig. Dann gibt es ein $\varepsilon > 0$ derart, dass zu jedem $n \geq 1$ Punkte $x_n, x'_n \in [a, b]$ existieren mit

$$|x_n - x'_n| < \frac{1}{n} \quad \text{und} \quad |f(x_n) - f(x'_n)| \geq \varepsilon.$$

Nach dem Satz von Bolzano-Weierstraß besitzt die beschränkte Folge (x_n) eine konvergente Teilfolge (x_{n_k}) . Für ihren Grenzwert gilt (nach §4, Corollar zu Satz 5)

$$\lim_{k \rightarrow \infty} x_{n_k} =: p \in [a, b].$$

Wegen $|x_{n_k} - x'_{n_k}| < \frac{1}{n_k}$ ist auch $\lim_{k \rightarrow \infty} x'_{n_k} = p$. Da f stetig ist, folgt daraus

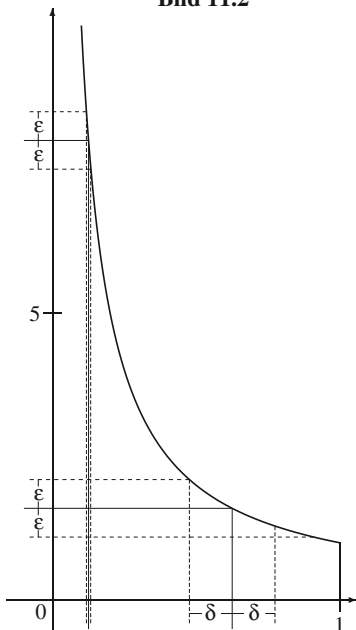
$$\lim_{k \rightarrow \infty} (f(x_{n_k}) - f(x'_{n_k})) = f(p) - f(p) = 0.$$

Dies ist ein Widerspruch zu $|f(x_{n_k}) - f(x'_{n_k})| \geq \varepsilon$ für alle k . Also ist die Annahme falsch und f gleichmäßig stetig.

Eine stetige Funktion auf einem nicht-kompakten Intervall ist jedoch i.Allg. nicht gleichmäßig stetig. Betrachten wir etwa die Funktion (siehe Bild 11.2)

$$f:]0, 1] \rightarrow \mathbb{R}, \quad x \mapsto \frac{1}{x}.$$

Bild 11.2



Diese Funktion ist natürlich stetig auf dem halboffenen Intervall $]0, 1]$. Hier lässt sich auch leicht explizit ein vom Punkt $p \in]0, 1]$ und $\varepsilon > 0$ abhängiges δ angeben. Wir setzen

$$\delta := \min\left(\frac{p}{2}, \frac{p^2\varepsilon}{2}\right).$$

Dann gilt für alle x mit $|x - p| < \delta$

$$\begin{aligned} |f(x) - f(p)| &= \left|\frac{1}{x} - \frac{1}{p}\right| = \left|\frac{x-p}{xp}\right| \\ &\leq \frac{2|x-p|}{p^2} < \frac{2\delta}{p^2} \leq \varepsilon, \end{aligned}$$

was die Stetigkeit von f im Punkt p zeigt. Das hier benutzte δ wird umso kleiner, je mehr man sich dem linken Rand des Intervalls nähert.

Die Funktion f ist im Intervall $]0, 1]$ nicht gleichmäßig stetig, da man nicht unabhängig von p wählen kann. Wäre f gleichmäßig stetig, gäbe es insbesondere zu $\varepsilon = 1$ ein $\delta > 0$, so dass

$$|f(x) - f(x')| < 1 \quad \text{für alle } x, x' \in]0, 1] \text{ mit } |x - x'| < \delta. \quad (*)$$

Es gibt aber ein $n \geq 1$ mit

$$\left|\frac{1}{n} - \frac{1}{2n}\right| < \delta \quad \text{und} \quad \left|f\left(\frac{1}{n}\right) - f\left(\frac{1}{2n}\right)\right| = n \geq 1,$$

was $(*)$ widerspricht.

Eine Folgerung aus der gleichmäßigen Stetigkeit ist der nächste Satz über die Approximierbarkeit stetiger Funktionen durch Treppenfunktionen, den wir später in der Integrationstheorie brauchen.

Satz 5. Sei $f: [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion. Dann gibt es zu jedem $\varepsilon > 0$ Treppenfunktionen $\varphi, \psi: [a, b] \rightarrow \mathbb{R}$ mit folgenden Eigenschaften:

- a) $\varphi(x) \leq f(x) \leq \psi(x)$ für alle $x \in [a, b]$,
- b) $|\varphi(x) - \psi(x)| \leq \varepsilon$ für alle $x \in [a, b]$.

Das Bild 11.3 veranschaulicht die Aussage von Satz 5.

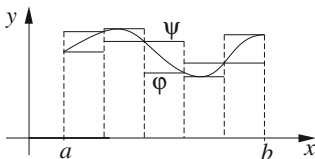


Bild 11.3

Beweis. Nach Satz 4 ist f gleichmäßig stetig. Zu $\varepsilon > 0$ existiert daher ein $\delta > 0$, so dass

$$|f(x) - f(x')| < \varepsilon \quad \text{für alle } x, x' \in [a, b] \text{ mit } |x - x'| < \delta.$$

Sei n so groß, dass $(b-a)/n < \delta$ und sei

$$t_k := a + k \frac{b-a}{n} \quad \text{für } k = 0, \dots, n.$$

Wir erhalten so eine (äquidistante) Intervallunterteilung

$$a = t_0 < t_1 < \dots < t_{n-1} < t_n = b.$$

mit $t_k - t_{k-1} < \delta$. Für $1 \leq k \leq n$ setzen wir

$$c_k := \sup\{f(x) : t_{k-1} \leq x \leq t_k\},$$

$$c'_k := \inf\{f(x) : t_{k-1} \leq x \leq t_k\}.$$

Da nach Satz 2 gilt $c_k = f(\xi_k)$ und $c'_k = f(\xi'_k)$ für gewisse Punkte $\xi_k, \xi'_k \in [t_{k-1}, t_k]$ und $|\xi_k - \xi'_k| < \delta$, folgt

$$|c_k - c'_k| < \varepsilon \quad \text{für alle } k.$$

Wir definieren nun Treppenfunktionen $\varphi, \psi: [a, b] \rightarrow \mathbb{R}$ wie folgt:

$$\varphi(a) := \psi(a) := f(a);$$

$$\varphi(x) := c'_k, \quad \psi(x) := c_k \quad \text{für } t_{k-1} < x \leq t_k, \quad (1 \leq k \leq n).$$

Damit sind die Bedingungen a) und b) erfüllt, q.e.d.

AUFGABEN

11.1. Es sei $F: [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion mit $F([a, b]) \subset [a, b]$. Man zeige, dass F mindestens einen Fixpunkt hat, d.h. es ein $x_0 \in [a, b]$ gibt mit $F(x_0) = x_0$.

11.2. Man zeige, dass die Funktion

$$\text{sqrt}: \mathbb{R}_+ \rightarrow \mathbb{R}, \quad x \mapsto \sqrt{x},$$

gleichmäßig stetig, die Funktion

$$\text{sq}: \mathbb{R}_+ \rightarrow \mathbb{R}, \quad x \mapsto x^2,$$

aber nicht gleichmäßig stetig ist.

11.3. Eine auf einer Teilmenge $D \subset \mathbb{R}$ definierte Funktion $f: D \rightarrow \mathbb{R}$ heißt *Lipschitz-stetig* mit Lipschitz-Konstante $L \in \mathbb{R}_+$ falls

$$|f(x) - f(x')| \leq L|x - x'| \quad \text{für alle } x, x' \in D.$$

a) Man zeige: Jede Lipschitz-stetige Funktion $f: D \rightarrow \mathbb{R}$ ist gleichmäßig stetig.

b) Die Funktion $f: [0, 1] \rightarrow \mathbb{R}, x \mapsto \sqrt{x}$ ist gleichmäßig stetig, aber nicht Lipschitz-stetig.

11.4. Sei $f: [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion. Der *Stetigkeitsmodul* $\omega_f: \mathbb{R}_+ \rightarrow \mathbb{R}$ von f ist wie folgt definiert:

$$\omega_f(\delta) := \sup \{ |f(x) - f(x')| : x, x' \in [a, b], |x - x'| \leq \delta \}.$$

Man beweise:

- i) ω_f ist stetig auf $\mathbb{R}_+$, insbesondere gilt $\lim_{\delta \searrow 0} \omega_f(\delta) = 0$.
- ii) Für $0 < \delta \leq \delta'$ gilt $\omega_f(\delta) \leq \omega_f(\delta')$.
- iii) Für alle $\delta, \delta' \in \mathbb{R}_+$ gilt $\omega_f(\delta + \delta') \leq \omega_f(\delta) + \omega_f(\delta')$.

11.5. Man beweise: Eine auf einem beschränkten offenen Intervall $]a, b[\subset \mathbb{R}$ stetige Funktion

$$f:]a, b[\longrightarrow \mathbb{R}$$

ist genau dann gleichmäßig stetig, wenn sie sich stetig auf das abgeschlossene Intervall $[a, b]$ fortsetzen lässt.

11.6. Man konstruiere ein Beispiel einer stetigen und beschränkten Funktion $f: [0, 1[\rightarrow \mathbb{R}$, die nicht gleichmäßig stetig ist.

11.7. Sei $f : \mathbb{R}_+ \rightarrow \mathbb{R}$ eine gleichmäßig stetige Funktion. Man zeige: Es gibt eine Konstante $M > 0$, so dass

$$|f(x)| \leq M(1+x) \quad \text{für alle } x \in \mathbb{R}_+.$$

11.8. Eine stetige Funktion $\varphi : [a, b] \rightarrow \mathbb{R}$ heißt stückweise linear, wenn es eine Unterteilung

$$a = t_0 < t_1 < \dots < t_r < t_{r+1} = b$$

des Intervalls $[a, b]$ und Konstanten α_k, β_k gibt, so dass für $k = 0, \dots, r$ gilt

$$\varphi(x) = \alpha_k + \beta_k x \quad \text{für } t_k \leq x \leq t_{k+1}.$$

(Der Graph von φ ist dann ein Polygonzug, der die Punkte $(t_k, \varphi(t_k))$, $k = 0, 1, \dots, r+1$ verbindet.)

Der Vektorraum aller stetigen, stückweise linearen Funktionen $\varphi : [a, b] \rightarrow \mathbb{R}$ werde mit $\text{PL}[a, b]$ bezeichnet (PL von *piecewise linear*).

Man zeige:

a) Jede Funktion $\varphi \in \text{PL}[a, b]$ lässt sich schreiben als

$$\varphi(x) = \alpha + \beta x + \sum_{k=1}^r c_k |x - t_k|$$

mit geeigneten Konstanten $\alpha, \beta, c_k \in \mathbb{R}$ und $t_j \in]a, b[$.

b) Jede stetige Funktion $f : [a, b] \rightarrow \mathbb{R}$ lässt sich gleichmäßig durch stetige, stückweise lineare Funktionen approximieren, d.h. zu jedem $\varepsilon > 0$ existiert eine Funktion $\varphi \in \text{PL}[a, b]$ mit

$$|f(x) - \varphi(x)| < \varepsilon \quad \text{für alle } x \in [a, b].$$

§ 12 Logarithmus und allgemeine Potenz

In diesem Paragraphen beweisen wir zunächst einen allgemeinen Satz über Umkehrfunktionen, den wir dann anwenden, um die Wurzeln und den Logarithmus zu definieren. Mithilfe des Logarithmus und der Exponentialfunktion wird dann die allgemeine Potenz a^x mit beliebiger positiver Basis a und reellem Exponenten x definiert.

Definition (Monotone Funktionen). Sei $D \subset \mathbb{R}$ und $f: D \rightarrow \mathbb{R}$ eine Funktion.

$$f \text{ hei\ss}t \left\{ \begin{array}{l} \text{monoton wachsend} \\ \text{streng monoton wachsend} \\ \text{monoton fallend} \\ \text{streng monoton fallend} \end{array} \right\}, \text{ falls } \left\{ \begin{array}{l} f(x) \leq f(x') \\ f(x) < f(x') \\ f(x) \geq f(x') \\ f(x) > f(x') \end{array} \right\}$$

für alle $x, x' \in D$ mit $x < x'$.

Satz 1. Sei $D \subset \mathbb{R}$ ein Intervall und $f: D \rightarrow \mathbb{R}$ eine stetige, streng monoton wachsende (oder fallende) Funktion. Dann bildet f das Intervall D bijektiv auf das Intervall $D' := f(D)$ ab, und die Umkehrfunktion

$$f^{-1}: D' \rightarrow \mathbb{R}$$

ist ebenfalls stetig und streng monoton wachsend (bzw. fallend).

Bemerkung. Die Umkehrfunktion ist genau genommen die Abbildung $f^{-1}: D' \rightarrow D$, definiert durch die Eigenschaft

$$f^{-1}(y) = x \iff f(x) = y.$$

Wir können aber f^{-1} unter Beibehaltung der Bezeichnung auch als Funktion $D' \rightarrow \mathbb{R}$ auffassen.

Vorsicht! Man verwechsle die Umkehrfunktion nicht mit der Funktion $x \mapsto 1/f(x)$.

Beweis von Satz 1. Wir haben bereits in §11 als Folgerung aus dem Zwischenwertsatz bewiesen, dass $D' = f(D)$ wieder ein Intervall ist. Als streng monotone Funktion ist f trivialerweise injektiv, bildet also D bijektiv auf D' ab, und die Umkehrabbildung ist wieder streng monoton (wachsend bzw. fallend). Es ist also nur noch die Stetigkeit von f^{-1} zu beweisen. Wir nehmen an, dass f streng monoton wächst (für streng monoton fallende Funktionen ist der Beweis analog zu führen). Sei $b \in D'$ ein gegebener Punkt und $a := f^{-1}(b)$, d.h.

$b = f(a)$. Wir zeigen, dass f^{-1} im Punkt b stetig ist. Wir behandeln zunächst den Fall, dass b weder rechter oder linker Randpunkt von D' ist, also auch a kein Randpunkt von D ist. Sei $\varepsilon > 0$ beliebig vorgegeben. Wir dürfen ohne Beschränkung der Allgemeinheit annehmen, dass ε so klein ist, dass das Intervall $[a - \varepsilon, a + \varepsilon]$ ganz in D liegt. Sei $b_1 := f(a - \varepsilon)$ und $b_2 := f(a + \varepsilon)$. Dann ist $b_1 < b < b_2$, und f bildet $[a - \varepsilon, a + \varepsilon]$ bijektiv auf das Intervall $[b_1, b_2]$ ab. Sei $\delta := \min(b - b_1, b_2 - b)$. Dann gilt

$$f^{-1}([b - \delta, b + \delta]) \subset]a - \varepsilon, a + \varepsilon[,$$

Dies zeigt (nach dem ε - δ -Kriterium), dass f^{-1} in b stetig ist. Ist $b \in D'$ rechter (bzw. linker) Randpunkt, so ist $a = f^{-1}(b)$ rechter (bzw. linker) Randpunkt von D und der Beweis verläuft ähnlich wie oben durch Betrachtung der Abbildung des Intervalls $[a - \varepsilon, a]$ (bzw. $[a, a + \varepsilon]$).

Wurzeln

Satz 2 und Definition. Sei k eine natürliche Zahl ≥ 2 . Die Funktion

$$f: \mathbb{R}_+ \longrightarrow \mathbb{R}, \quad x \mapsto x^k,$$

ist streng monoton wachsend und bildet $\mathbb{R}_+$ bijektiv auf $\mathbb{R}_+$ ab. Die Umkehrfunktion

$$f^{-1}: \mathbb{R}_+ \longrightarrow \mathbb{R}, \quad x \mapsto \sqrt[k]{x},$$

ist stetig und streng monoton wachsend und wird als k -te Wurzel bezeichnet.

Beweis. Es ist klar, dass f streng monoton wächst und das Intervall $[0, +\infty[$ stetig und bijektiv auf $[0, +\infty[$ abbildet. Somit folgt Satz 2 unmittelbar aus Satz 1.

Bemerkung. Falls k ungerade ist, ist die Funktion

$$f: \mathbb{R} \longrightarrow \mathbb{R}, \quad x \mapsto x^k,$$

streng monoton und bijektiv. In diesem Fall kann also die k -te Wurzel als Funktion

$$\mathbb{R} \longrightarrow \mathbb{R}, \quad x \mapsto \sqrt[k]{x},$$

auf ganz $\mathbb{R}$ definiert werden.

Natürlicher Logarithmus

Satz 3 und Definition. Die Exponentialfunktion $\exp: \mathbb{R} \rightarrow \mathbb{R}$ ist streng monoton wachsend und bildet $\mathbb{R}$ bijektiv auf $\mathbb{R}_+^*$ ab. Die Umkehrfunktion

$$\log: \mathbb{R}_+^* \longrightarrow \mathbb{R}$$

ist stetig und streng monoton wachsend und heißt natürlicher Logarithmus (Bild 12.1). Es gilt die Funktionalgleichung

$$\log(xy) = \log x + \log y \quad \text{für alle } x, y \in \mathbb{R}_+^*.$$

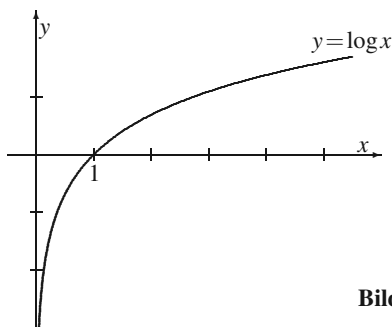


Bild 12.1 Logarithmus

Bemerkung. Statt $\log$ ist auch (wie in früheren Auflagen dieses Buches) die Bezeichnung $\ln$ gebräuchlich.

Beweis. a) Wir zeigen zunächst, dass die Funktion $\exp$ streng monoton wächst. Für $\xi > 0$ gilt

$$\exp(\xi) = 1 + \xi + \frac{\xi^2}{2} + \dots > 1.$$

Sei $x < x'$. Dann ist $\xi := x' - x > 0$, also $\exp(\xi) > 1$. Daraus folgt

$$\exp(x') = \exp(x + \xi) = \exp(x) \exp(\xi) > \exp(x),$$

d.h. $\exp$ ist streng monoton wachsend.

b) Für alle $n \in \mathbb{N}$ gilt

$$\exp(n) \geq 1 + n$$

und

$$\exp(-n) = \frac{1}{\exp(n)} \leq \frac{1}{1+n}.$$

Daraus folgt

$$\lim_{n \rightarrow \infty} \exp(n) = \infty \quad \text{und} \quad \lim_{n \rightarrow \infty} \exp(-n) = 0.$$

Also gilt $\exp(\mathbb{R}) =]0, \infty[= \mathbb{R}_+^*$ und nach Satz 1 ist die Umkehrfunktion $\log: \mathbb{R}_+^* \rightarrow \mathbb{R}$ stetig und streng monoton wachsend.

c) Zum Beweis der Funktionalgleichung setzen wir $\xi := \log(x)$ und $\eta := \log(y)$. Dann ist nach Definition $\exp(\xi) = x$ und $\exp(\eta) = y$. Aus der Funktionalgleichung der Exponentialfunktion folgt

$$\exp(\xi + \eta) = \exp(\xi) \exp(\eta) = xy.$$

Wieder nach Definition der Umkehrfunktion ist daher

$$\log(xy) = \xi + \eta = \log(x) + \log(y), \quad \text{q.e.d.}$$

Definition (Exponentialfunktion zur Basis a). Für $a > 0$ sei die Funktion $\exp_a: \mathbb{R} \rightarrow \mathbb{R}$ definiert durch

$$\exp_a(x) := \exp(x \log a).$$

Satz 4. Die Funktion $\exp_a: \mathbb{R} \rightarrow \mathbb{R}$ ist stetig und es gilt:

- i) $\exp_a(x+y) = \exp_a(x) \exp_a(y)$ für alle $x, y \in \mathbb{R}$.
- ii) $\exp_a(n) = a^n$ für alle $n \in \mathbb{Z}$.
- iii) $\exp_a(\frac{p}{q}) = \sqrt[q]{a^p}$ für alle $p \in \mathbb{Z}$ und $q \in \mathbb{N}$ mit $q \geq 2$.

Beweis. a) Die Funktion $\exp_a$ ist die Komposition der stetigen Funktionen $x \mapsto x \log a$ und $y \mapsto \exp(y)$, also nach §10, Satz 2, selbst stetig.

b) Die Behauptung i) folgt unmittelbar aus der Funktionalgleichung der Exponentialfunktion. Aus i) ergibt sich, wenn man $y = -x$ setzt, insbesondere

$$\exp_a(-x) = \frac{1}{\exp_a(x)}.$$

c) Durch vollständige Induktion zeigt man

$$\exp_a(nx) = (\exp_a(x))^n \quad \text{für alle } n \in \mathbb{N} \text{ und } x \in \mathbb{R}.$$

Da $\exp_a(1) = \exp(\log a) = a$ und $\exp_a(-1) = 1/a$, folgt daraus mit $x = 1$ bzw. $x = -1$

$$\exp_a(n) = a^n \quad \text{und} \quad \exp_a(-n) = a^{-n}.$$

Damit ist ii) bewiesen. Weiter ergibt sich

$$a^p = \exp_a(p) = \exp_a\left(q \cdot \frac{p}{q}\right) = \left(\exp_a\left(\frac{p}{q}\right)\right)^q,$$

also durch Ziehen der q -ten Wurzel die Behauptung iii).

Corollar. Für alle $a > 0$ gilt $\lim_{n \rightarrow \infty} \sqrt[n]{a} = 1$.

Beweis. Dies folgt aus der Stetigkeit der Funktion $\exp_a$:

$$\lim_{n \rightarrow \infty} \sqrt[n]{a} = \lim_{n \rightarrow \infty} \exp_a\left(\frac{1}{n}\right) = \exp_a(0) = 1.$$

Bezeichnung. Satz 4 rechtfertigt die Bezeichnung

$$a^x := \exp_a(x) = \exp(x \log a).$$

Da $\log e = 1$, ist insbesondere $e^x = \exp(x) = \exp_e(x)$.

Die für ganzzahlige Potenzen bekannten Rechenregeln gelten auch für die allgemeine Potenz.

Satz 5 (Rechenregeln für Potenzen). Für alle $a, b \in \mathbb{R}_+^*$ und $x, y \in \mathbb{R}$ gilt:

- i) $a^x a^y = a^{x+y}$,
- ii) $(a^x)^y = a^{xy}$,
- iii) $a^x b^x = (ab)^x$,
- iv) $(1/a)^x = a^{-x}$.

Beweis. Die Regel i) ist nur eine andere Schreibweise von Satz 4 i).

Zu ii) Da $a^x = \exp(x \log a)$, ist $\log(a^x) = x \log a$, also

$$(a^x)^y = \exp(y \log(a^x)) = \exp(yx \log a) = a^{xy}.$$

Die Behauptungen iii) und iv) sind ebenso einfach zu beweisen.

Wir zeigen jetzt, dass die Funktionalgleichung $a^{x+y} = a^x a^y$ charakteristisch für die allgemeine Potenz ist.

Satz 6. Sei $F: \mathbb{R} \rightarrow \mathbb{R}$ eine stetige Funktion mit

$$F(x+y) = F(x)F(y) \quad \text{für alle } x, y \in \mathbb{R}.$$

Dann ist entweder $F(x) = 0$ für alle $x \in \mathbb{R}$ oder es ist $a := F(1) > 0$ und

$$F(x) = a^x \quad \text{für alle } x \in \mathbb{R}.$$

Beweis. Da $F(1) = F(\frac{1}{2})^2$, gilt in jedem Fall $F(1) \geq 0$.

a) Setzen wir zunächst voraus, dass $a := F(1) > 0$. Da

$$a = F(1+0) = F(1)F(0) = aF(0),$$

folgt daraus $F(0) = 1$. Man beweist nun wie in Satz 4 allein mithilfe der Funktionalgleichung

$$\begin{aligned} F(n) &= a^n && \text{für alle } n \in \mathbb{Z}, \\ F\left(\frac{p}{q}\right) &= \sqrt[q]{a^p} && \text{für alle } p \in \mathbb{Z} \text{ und } q \in \mathbb{N} \text{ mit } q \geq 2. \end{aligned}$$

Es gilt also $F(x) = a^x$ für alle rationalen Zahlen x . Sei nun x eine beliebige reelle Zahl. Dann gibt es eine Folge $(x_n)_{n \in \mathbb{N}}$ rationaler Zahlen mit $\lim_{n \rightarrow \infty} x_n = x$. Wegen der Stetigkeit der Funktionen F und $\exp_a$ folgt daraus

$$F(x) = \lim_{n \rightarrow \infty} F(x_n) = \lim_{n \rightarrow \infty} a^{x_n} = a^x.$$

b) Es bleibt noch der Fall $F(1) = 0$ zu untersuchen. Wir haben zu zeigen, dass dann $F(x) = 0$ für alle $x \in \mathbb{R}$. Dies sieht man so:

$$F(x) = F(1 + (x-1)) = F(1)F(x-1) = 0 \cdot F(x-1) = 0, \quad \text{q.e.d.}$$

Bemerkung. Die Definition $a^x := \exp(x \log a)$ mag zunächst künstlich erscheinen. Wenn man aber die Definition so treffen will, dass $a^{x+y} = a^x a^y$ für alle $x, y \in \mathbb{R}$ sowie $a^1 = a$, und dass a^x stetig von x abhängt, so sagt Satz 6, dass notwendig $a^x = \exp(x \log a)$ ist.

Berechnung einiger Grenzwerte

Wir beweisen jetzt einige wichtige Aussagen über das Verhalten des Logarithmus und der Potenzfunktionen für $x \rightarrow \infty$ und $x \rightarrow 0$.

(12.1) Für alle $k \in \mathbb{N}$ gilt $\lim_{x \rightarrow \infty} \frac{e^x}{x^k} = \infty$.

Man drückt dies auch so aus: e^x wächst für $x \rightarrow \infty$ schneller gegen unendlich, als jede Potenz von x .

Beweis. Für alle $x > 0$ ist

$$e^x = \sum_{n=0}^{\infty} \frac{x^n}{n!} > \frac{x^{k+1}}{(k+1)!},$$

also $\frac{e^x}{x^k} > \frac{x}{(k+1)!}$. Daraus folgt die Behauptung.

(12.2) Für alle $k \in \mathbb{N}$ gilt

$$\lim_{x \rightarrow \infty} x^k e^{-x} = 0 \quad \text{und} \quad \lim_{x \searrow 0} x^k e^{1/x} = \infty.$$

Beweis. Die erste Aussage folgt aus (12.1), da $x^k e^{-x} = \left(\frac{e^x}{x^k}\right)^{-1}$. Die zweite Aussage folgt ebenfalls aus (12.1), denn

$$\lim_{x \searrow 0} x^k e^{1/x} = \lim_{y \rightarrow \infty} \left(\frac{1}{y}\right)^k e^y = \lim_{y \rightarrow \infty} \frac{e^y}{y^k} = \infty.$$

(12.3) $\lim_{x \rightarrow \infty} \log x = \infty$ und $\lim_{x \searrow 0} \log x = -\infty$.

Beweis. Sei $K \in \mathbb{R}$ beliebig vorgegeben. Da die Funktion $\log$ streng monoton wächst, gilt $\log x > K$ für alle $x > e^K$. Also ist $\lim_{x \rightarrow \infty} \log x = \infty$. Daraus folgt die zweite Behauptung, da

$$\lim_{x \searrow 0} \log x = \lim_{y \rightarrow \infty} \log(1/y) = -\lim_{y \rightarrow \infty} \log y = -\infty.$$

(12.4) Für jede reelle Zahl $\alpha > 0$ gilt

$$\lim_{x \searrow 0} x^\alpha = 0 \quad \text{und} \quad \lim_{x \searrow 0} x^{-\alpha} = \infty.$$

Beweis. Sei $(x_n)_{n \in \mathbb{N}}$ eine Folge reeller Zahlen mit $x_n > 0$ und $\lim_{n \rightarrow \infty} x_n = 0$. Mit (12.3) folgt

$$\lim_{n \rightarrow \infty} \alpha \log x_n = -\infty.$$

Da nach (12.2) gilt $\lim_{y \rightarrow -\infty} e^y = 0$, folgt

$$\lim_{n \rightarrow \infty} x_n^\alpha = \lim_{n \rightarrow \infty} e^{\alpha \log x_n} = 0,$$

also $\lim_{x \searrow 0} x^\alpha = 0$. Die zweite Behauptung gilt wegen $x^{-\alpha} = \frac{1}{x^\alpha}$.

Bemerkung. Wegen (12.4) definiert man

$$0^\alpha := 0 \quad \text{für alle } \alpha > 0.$$

Man erhält dann eine auf ganz $\mathbb{R}_+ = [0, \infty[$ stetige Funktion

$$\mathbb{R}_+ \longrightarrow \mathbb{R}, \quad x \mapsto x^\alpha.$$

(12.5) Für alle $\alpha > 0$ gilt $\lim_{x \rightarrow \infty} \frac{\log x}{x^\alpha} = 0$.

Anders ausgedrückt: Der Logarithmus wächst für $x \rightarrow \infty$ langsamer gegen unendlich, als jede positive Potenz von x .

Beweis. Sei (x_n) eine Folge positiver Zahlen mit $\lim x_n = \infty$. Für die Folge $y_n := \alpha \log x_n$ gilt wegen (12.3) dann ebenfalls $\lim y_n = \infty$. Da $x_n^\alpha = e^{y_n}$, erhalten wir unter Verwendung von (12.2)

$$\lim_{n \rightarrow \infty} \frac{\log x_n}{x_n^\alpha} = \lim_{n \rightarrow \infty} \frac{1}{\alpha} y_n e^{-y_n} = 0, \quad \text{q.e.d.}$$

(12.6) Für alle $\alpha > 0$ gilt $\lim_{x \searrow 0} x^\alpha \log x = 0$.

Dies folgt aus (12.5), da $x^\alpha \log x = -\frac{\log(1/x)}{(1/x)^\alpha}$.

$$(12.7) \quad \lim_{\substack{x \rightarrow 0 \\ x \neq 0}} \frac{e^x - 1}{x} = 1.$$

Beweis. Nach §8, Satz 2, gilt

$$|e^x - (1+x)| \leq |x|^2 \quad \text{für } |x| \leq \frac{3}{2}.$$

Division durch $|x|$ ergibt für $0 < |x| \leq \frac{3}{2}$

$$\left| \frac{e^x - 1}{x} - 1 \right| = \left| \frac{e^x - (1+x)}{x} \right| \leq |x|.$$

Daraus folgt die Behauptung.

Die Landau-Symbole O und o

E. Landau hat zum Vergleich des Wachstums von Funktionen suggestive Bezeichnungen eingeführt, die wir jetzt vorstellen. Gegeben seien zwei Funktionen

$$f, g:]a, \infty[\longrightarrow \mathbb{R}.$$

Dann schreibt man

$$f(x) = o(g(x)) \quad \text{für } x \rightarrow \infty,$$

(gesprochen: $f(x)$ gleich klein-oh von $g(x)$), wenn zu jedem $\varepsilon > 0$ ein $R > a$ existiert, so dass

$$|f(x)| \leq \varepsilon |g(x)| \quad \text{für alle } x \geq R.$$

Ist $g(x) \neq 0$ für $x \geq R_0$, so ist dies äquivalent zu

$$\lim_{x \rightarrow \infty} \frac{f(x)}{g(x)} = 0.$$

Die Bedingung $f(x) = o(g(x))$ sagt also anschaulich, dass f asymptotisch für $x \rightarrow \infty$ im Vergleich zu g verschwindend klein ist. Damit lässt sich z.B. nach (12.2) und (12.5) schreiben

$$e^{-x} = o(x^{-n}) \quad \text{für } x \rightarrow \infty$$

für alle $n \in \mathbb{N}$ und

$$\log x = o(x^\alpha), \quad (\alpha > 0, x \rightarrow \infty).$$

Man beachte jedoch, dass das Gleichheitszeichen in $f(x) = o(g(x))$ nicht eine Gleichheit von Funktionen bedeutet, sondern nur eine Eigenschaft der Funktion f im Vergleich zu g ausdrückt. So folgt natürlich aus $f_1(x) = o(g(x))$ und $f_2(x) = o(g(x))$ nicht, dass $f_1 = f_2$, aber z.B.

$$f_1(x) - f_2(x) = o(g(x)) \quad \text{und} \quad f_1(x) + f_2(x) = o(g(x)).$$

Das Symbol O ist für zwei Funktionen $f, g:]a, \infty[\rightarrow \mathbb{R}$ so definiert: Man schreibt

$$f(x) = O(g(x)) \quad \text{für } x \rightarrow \infty,$$

wenn Konstanten $K \in \mathbb{R}_+$ und $R > a$ existieren, so dass

$$|f(x)| \leq K|g(x)| \quad \text{für alle } x \geq R.$$

Falls $g(x) \neq 0$ für $x \geq R_0$, ist dies äquivalent mit

$$\limsup_{x \rightarrow \infty} \left| \frac{f(x)}{g(x)} \right| < \infty.$$

Anschaulich bedeutet das, dass asymptotisch für $x \rightarrow \infty$ die Funktion f höchstens von gleicher Größenordnung wie g ist. Z.B. gilt für jedes Polynom n -ten Grades

$$P(x) = a_0 + a_1x + \dots + a_{n-1}x^{n-1} + a_nx^n,$$

dass $P(x) = O(x^n)$ für $x \rightarrow \infty$.

Die Landau-Symbole o und O sind nicht nur für den Grenzübergang $x \rightarrow \infty$, sondern auch für andere Grenzübergänge $x \rightarrow x_0$ definiert. Seien etwa $f, g: D \rightarrow \mathbb{R}$ zwei auf der Teilmenge $D \subset \mathbb{R}$ definierte Funktionen und x_0 ein Berührungspunkt von D . Dann schreibt man

$$f(x) = o(g(x)) \quad \text{für } x \rightarrow x_0, x \in D,$$

falls zu jedem $\varepsilon > 0$ ein $\delta > 0$ existiert, so dass

$$|f(x)| \leq \varepsilon |g(x)| \quad \text{für alle } x \in D \text{ mit } |x - x_0| < \delta.$$

Falls $g(x) \neq 0$ in D , ist dies wieder gleichbedeutend mit

$$\lim_{\substack{x \rightarrow x_0 \\ x \in D}} \frac{f(x)}{g(x)} = 0.$$

Damit schreibt sich (12.6) als

$$\log x = o\left(\frac{1}{x^\alpha}\right) \quad (\alpha > 0, x \searrow 0),$$

und aus (12.2) folgt für alle $n \in \mathbb{N}$

$$e^{-1/x} = o(x^n) \quad \text{für } x \searrow 0.$$

Manchmal ist folgende Erweiterung der Schreibweise nützlich:

$$f_1(x) = f_2(x) + o(g(x)) \quad \text{für } x \rightarrow x_0$$

bedeute $f_1(x) - f_2(x) = o(g(x))$. Sei beispielsweise $f: D \rightarrow \mathbb{R}$ eine Funktion und $x_0 \in D$. Dann ist

$$f(x) = f(x_0) + o(1) \quad \text{für } x \rightarrow x_0$$

gleichbedeutend mit $\lim_{x \rightarrow x_0} (f(x) - f(x_0)) = 0$, also mit der Stetigkeit von f in x_0 . Die Aussage von Beispiel (12.8) ist äquivalent zu

$$\log(1+x) = x + o(x) \quad \text{für } x \rightarrow 0.$$

Analoge Schreibweisen führt man für das Symbol O ein. Z.B. gilt, vgl. (12.7),

$$e^x = 1 + x + O(x^2) \quad \text{für } x \rightarrow 0.$$

AUFGABEN

12.1. Man zeige: Die Funktion

$$\exp_a: \mathbb{R} \longrightarrow \mathbb{R}, \quad x \mapsto a^x,$$

ist für $a > 1$ streng monoton wachsend und für $0 < a < 1$ streng monoton fallend. In beiden Fällen wird $\mathbb{R}$ bijektiv auf $\mathbb{R}_+^*$ abgebildet. Die Umkehrfunktion

$${}^a\log: \mathbb{R}_+^* \longrightarrow \mathbb{R}$$

(Logarithmus zur Basis a) ist stetig und es gilt

$${}^a\log x = \frac{\log x}{\log a} \quad \text{für alle } x \in \mathbb{R}_+^*.$$

12.2. Man zeige: Die Funktion $\sinh$ bildet $\mathbb{R}$ bijektiv auf $\mathbb{R}$ ab; die Funktion $\cosh$ bildet $\mathbb{R}_+$ bijektiv auf $[1, \infty[$ ab.

(Die Funktionen $\sinh$ und $\cosh$ wurden in Aufgabe 10.1 definiert.)

Für die Umkehrfunktionen

$$\operatorname{Arsinh} : \mathbb{R} \longrightarrow \mathbb{R} \quad (\text{Area sinus hyperbolici}),$$

$$\operatorname{Arcosh} : [1, \infty[\longrightarrow \mathbb{R} \quad (\text{Area cosinus hyperbolici})$$

gelten die Beziehungen

$$\operatorname{Arsinh} x = \log(x + \sqrt{x^2 + 1}),$$

$$\operatorname{Arcosh} x = \log(x + \sqrt{x^2 - 1}).$$

12.3. Sei $D \subset \mathbb{R}$ ein Intervall und $f: D \longrightarrow \mathbb{R}$ eine streng monotone Funktion (nicht notwendig stetig). Sei $D' := f(D)$. Man beweise: Die Umkehrfunktion $f^{-1}: D' \longrightarrow D \subset \mathbb{R}$ ist stetig.

12.4. Man beweise:

$$\lim_{x \searrow 0} x^x = 1 \quad \text{und} \quad \lim_{n \rightarrow \infty} \sqrt[n]{n} = 1.$$

12.5. Sei $a > 0$. Die Folgen (x_n) und (y_n) seien definiert durch

$$x_0 := a, \quad x_{n+1} := \sqrt{x_n},$$

$$y_n := 2^n(x_n - 1).$$

Man beweise $\lim_{n \rightarrow \infty} y_n = \log a$.

Hinweis. Man verwende $\lim_{x \rightarrow 0} \frac{e^x - 1}{x} = 1$.

12.6. Man beweise

$$\lim_{\substack{x \rightarrow 0 \\ x \neq 0}} \frac{\log(1+x)}{x} = 1.$$

12.7. Man zeige:

- i) $\log(1+x) \leq x$ für alle $x \geq 0$.
- ii) $\log(1-x) \geq -2x$ für $0 \leq x \leq \frac{1}{2}$.
- iii) $|\log(1+x)| \leq 2|x|$ für $|x| \leq \frac{1}{2}$.

12.8. Man zeige: Für alle $n \in \mathbb{N}$ und alle $\alpha > 0$ gilt für $x \rightarrow \infty$:

- i) $x(\log x)^n = o(x^{1+\alpha})$,
- ii) $x^n = o(e^{\sqrt{x}})$,
- iii) $e^{\sqrt{x}} = o(e^{\alpha x})$.

12.9. Man bestimme alle stetigen Funktionen, die folgenden Funktionalgleichungen genügen:

- i) $f: \mathbb{R} \rightarrow \mathbb{R}, \quad f(x+y) = f(x) + f(y)$,
- ii) $g: \mathbb{R}_+^* \rightarrow \mathbb{R}, \quad g(xy) = g(x) + g(y)$,
- iii) $h: \mathbb{R}_+^* \rightarrow \mathbb{R}, \quad h(xy) = h(x)h(y)$.

12.10. Seien $f_1, f_2, g_1, g_2:]a, \infty[\rightarrow \mathbb{R}$ Funktionen mit

$$f_1(x) = o(g_1(x)) \quad \text{und} \quad f_2(x) = O(g_2(x)) \quad \text{für } x \rightarrow \infty.$$

Man zeige $f_1(x)f_2(x) = o(g_1(x)g_2(x))$ für $x \rightarrow \infty$.

12.11. Seien $f, g:]-\varepsilon, \varepsilon[\rightarrow \mathbb{R}$, ($\varepsilon > 0$), Funktionen mit

$$\left. \begin{aligned} f(x) &= a_0 + a_1x + a_2x^2 + \dots + a_nx^n + o(|x|^n) \\ g(x) &= b_0 + b_1x + b_2x^2 + \dots + b_nx^n + o(|x|^n) \end{aligned} \right\} \quad \text{für } x \rightarrow 0$$

Man zeige

$$f(x)g(x) = c_0 + c_1x + c_2x^2 + \dots + c_nx^n + o(|x|^n) \quad \text{für } x \rightarrow 0,$$

wobei $c_k = \sum_{i=0}^k a_i b_{k-i}$.

12.12. Sei $f:]-\varepsilon, \varepsilon[\rightarrow \mathbb{R}$, ($\varepsilon > 0$), eine Funktion mit

$$f(x) = 1 + ax + O(x^2) \quad \text{für } x \rightarrow 0.$$

Man zeige:

$$\frac{1}{f(x)} = 1 - ax + O(x^2) \quad \text{für } x \rightarrow 0.$$

§ 13 Die Exponentialfunktion im Komplexen

Wir wollen im nächsten Paragraphen die trigonometrischen Funktionen vermöge der Eulerschen Formel $e^{ix} = \cos x + i \sin x$ einführen. Zu diesem Zweck brauchen wir die Exponentialfunktion für komplexe Argumente. Sie ist wie im Reellen durch die Exponentialreihe definiert. Dazu müssen wir einige Sätze über die Konvergenz von Folgen und Reihen ins Komplexe übertragen, was eine gute Gelegenheit zur Wiederholung dieser Begriffe gibt.

Der Körper der komplexen Zahlen

Die Menge $\mathbb{R} \times \mathbb{R}$ aller (geordneten) Paare reeller Zahlen bildet zusammen mit der Addition und Multiplikation

$$\begin{aligned}(x_1, y_1) + (x_2, y_2) &:= (x_1 + x_2, y_1 + y_2), \\ (x_1, y_1) \cdot (x_2, y_2) &:= (x_1 x_2 - y_1 y_2, x_1 y_2 + y_1 x_2),\end{aligned}$$

einen Körper. Das Nullelement ist $(0, 0)$, das Einselement $(1, 0)$. Das Inverse eines Elements $(x, y) \neq (0, 0)$ ist

$$(x, y)^{-1} = \left(\frac{x}{x^2 + y^2}, \frac{-y}{x^2 + y^2} \right).$$

Man prüft leicht alle Körper-Axiome nach. Nur die Verifikation des Assoziativ-Gesetzes der Multiplikation und des Distributiv-Gesetzes erfordert eine etwas längere (aber einfache) Rechnung. Wir führen dies für das Assoziativ-Gesetz durch:

$$\begin{aligned}((x_1, y_1)(x_2, y_2))(x_3, y_3) &= (x_1 x_2 - y_1 y_2, x_1 y_2 + y_1 x_2)(x_3, y_3) = \\ (x_1 x_2 - y_1 y_2)x_3 - (x_1 y_2 + y_1 x_2)y_3, & (x_1 x_2 - y_1 y_2)y_3 + (x_1 y_2 + y_1 x_2)x_3.\end{aligned}$$

Andrerseits ist

$$\begin{aligned}(x_1, y_1)((x_2, y_2)(x_3, y_3)) &= (x_1, y_1)(x_2 x_3 - y_2 y_3, x_2 y_3 + y_2 x_3) = \\ (x_1(x_2 x_3 - y_2 y_3) - y_1(x_2 y_3 + y_2 x_3), & x_1(x_2 y_3 + y_2 x_3) + y_1(x_2 x_3 - y_2 y_3)).\end{aligned}$$

Aufgrund des Assoziativ- und Distributiv-Gesetzes für den Körper $\mathbb{R}$ sieht man, dass beide Ausdrücke gleich sind.

Der entstandene Körper heißt Körper der komplexen Zahlen und wird mit $\mathbb{C}$ bezeichnet. Für die speziellen komplexen Zahlen der Gestalt $(x, 0)$ gilt

$$\begin{aligned}(x_1, 0) + (x_2, 0) &= (x_1 + x_2, 0), \\ (x_1, 0) \cdot (x_2, 0) &= (x_1 x_2, 0),\end{aligned}$$

sie werden also genau so wie die entsprechenden reellen Zahlen addiert und multipliziert; wir dürfen deshalb die reelle Zahl x mit der komplexen Zahl $(x, 0)$ identifizieren. $\mathbb{R}$ wird so eine Teilmenge von $\mathbb{C}$. Eine wichtige komplexe Zahl ist die sog. *imaginäre Einheit* $i := (0, 1)$; für sie gilt

$$i^2 = (0, 1)(0, 1) = (-1, 0) = -1,$$

sie löst also die Gleichung $z^2 + 1 = 0$ in $\mathbb{C}$. Mithilfe von i erhält man die gebräuchliche Schreibweise für die komplexen Zahlen

$$z = (x, y) = (x, 0) + (0, 1)(y, 0) = x + iy, \quad x, y \in \mathbb{R}.$$

Man veranschaulicht sich die komplexen Zahlen in der *Gauß'schen Zahlenebene* (Bild 13.1). Die Addition zweier komplexer Zahlen wird dann die gewöhnliche Vektoraddition (Bild 13.2). Eine geometrische Deutung der Multiplikation werden wir im nächsten Paragraphen kennenlernen.

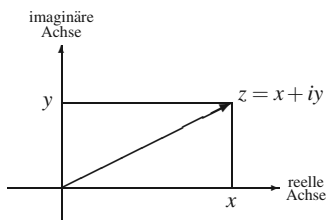


Bild 13.1

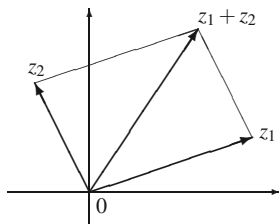


Bild 13.2

Für eine komplexe Zahl $z = x + iy$, ($x, y \in \mathbb{R}$), werden *Realteil* und *Imaginärteil* wie folgt definiert:

$$\operatorname{Re}(z) := x, \quad \text{und} \quad \operatorname{Im}(z) := y.$$

Zwei komplexe Zahlen z, z' sind also genau dann gleich, wenn

$$\operatorname{Re}(z) = \operatorname{Re}(z') \quad \text{und} \quad \operatorname{Im}(z) = \operatorname{Im}(z').$$

Komplexe Konjugation

Für eine komplexe Zahl $z = x + iy$, ($x, y \in \mathbb{R}$), definiert man die konjugiert komplexe Zahl durch

$$\bar{z} := x - iy.$$

In der Gauß'schen Zahlenebene entsteht $\bar{z}$ aus z durch Spiegelung an der reellen Achse. Offenbar gilt $\bar{z} = z$ genau dann, wenn z reell ist. Aus der Definition folgt

$$\operatorname{Re}(z) = \frac{1}{2}(z + \bar{z}), \quad \operatorname{Im}(z) = \frac{1}{2i}(z - \bar{z}).$$

Einfach nachzurechnen sind folgende Rechenregeln für die Konjugation: Für alle $z, w \in \mathbb{C}$ gilt

- a) $\overline{\bar{z}} = z$,
- b) $\overline{z + w} = \bar{z} + \bar{w}$,
- c) $\overline{z \cdot w} = \bar{z} \cdot \bar{w}$.

Die Regeln a) – c) besagen, dass die Abbildung $z \mapsto \bar{z}$ ein *involutorischer Automorphismus* von $\mathbb{C}$ ist.

Betrag einer komplexen Zahl. Sei $z = x + iy \in \mathbb{C}$. Dann ist

$$z\bar{z} = (x + iy)(x - iy) = x^2 + y^2$$

eine nicht-negative reelle Zahl. Man setzt

$$|z| := \sqrt{z\bar{z}} \in \mathbb{R}_+.$$

$|z|$ heißt der Betrag von z . Da $|z| = \sqrt{x^2 + y^2}$, ist der Betrag von z gleich dem Abstand des Punktes z vom Nullpunkt der Gauß'schen Zahlenebene bzgl. der gewöhnlichen euklidischen Metrik.

Für $z \in \mathbb{R}$ stimmt der Betrag mit dem Betrag für reelle Zahlen überein. Für alle $z \in \mathbb{C}$ gilt $|z| = |\bar{z}|$.

Satz 1. *Der Betrag in $\mathbb{C}$ hat folgende Eigenschaften:*

- a) *Es ist $|z| \geq 0$ für alle $z \in \mathbb{C}$ und*

$$|z| = 0 \iff z = 0.$$

- b) (Multiplikativität)

$$|z_1 z_2| = |z_1| \cdot |z_2| \quad \text{für alle } z_1, z_2 \in \mathbb{C}.$$

- c) (Dreiecks-Ungleichung)

$$|z_1 + z_2| \leq |z_1| + |z_2| \quad \text{für alle } z_1, z_2 \in \mathbb{C}.$$

Bemerkungen. Satz 1 sagt, dass $\mathbb{C}$ durch den Betrag $z \mapsto |z|$ zu einem bewerteten Körper wird, vgl. die Bemerkung zu §3, Satz 1.

Die Dreiecks-Ungleichung drückt aus, dass in dem Dreieck mit den Ecken $0, z_1, z_1 + z_2$ (vgl. Bild 13.2) die Länge der Seite von 0 nach $z_1 + z_2$ kleiner-gleich der Summe der Längen der beiden anderen Seiten ist.

Beweis von Satz 1. Die Behauptung a) ist trivial.

Zu b) Nach Definition des Betrages ist

$$|z_1 z_2|^2 = (z_1 z_2)(\overline{z_1 z_2}) = z_1 z_2 \bar{z}_1 \bar{z}_2 = (z_1 \bar{z}_1)(z_2 \bar{z}_2) = |z_1|^2 |z_2|^2.$$

Indem man die Wurzel zieht, erhält man die Behauptung.

Zu c) Da für jede komplexe Zahl gilt $\operatorname{Re}(z) \leq |z|$, folgt

$$\operatorname{Re}(z_1 \bar{z}_2) \leq |z_1 \bar{z}_2| = |z_1| |\bar{z}_2| = |z_1| |z_2|.$$

Nun ist

$$\begin{aligned} |z_1 + z_2|^2 &= (z_1 + z_2)(\bar{z}_1 + \bar{z}_2) = z_1 \bar{z}_1 + z_1 \bar{z}_2 + z_2 \bar{z}_1 + z_2 \bar{z}_2 \\ &= |z_1|^2 + 2 \operatorname{Re}(z_1 \bar{z}_2) + |z_2|^2 \\ &\leq |z_1|^2 + 2|z_1||z_2| + |z_2|^2 = (|z_1| + |z_2|)^2, \end{aligned}$$

also $|z_1 + z_2| \leq |z_1| + |z_2|$, q.e.d.

Konvergenz in $\mathbb{C}$

Wir übertragen nun die wichtigsten Begriffe und Sätze aus §4, §5, §7 über Konvergenz auf Folgen und Reihen komplexer Zahlen.

Definition. Eine Folge $(c_n)_{n \in \mathbb{N}}$ komplexer Zahlen heißt *konvergent* gegen eine komplexe Zahl c , falls zu jedem $\varepsilon > 0$ ein $N \in \mathbb{N}$ existiert, so dass

$$|c_n - c| < \varepsilon \quad \text{für alle } n \geq N.$$

Wir schreiben dann $\lim_{n \rightarrow \infty} c_n = c$.

Satz 2. Sei $(c_n)_{n \in \mathbb{N}}$ eine Folge komplexer Zahlen. Die Folge konvergiert genau dann, wenn die beiden reellen Folgen $(\operatorname{Re}(c_n))_{n \in \mathbb{N}}$ und $(\operatorname{Im}(c_n))_{n \in \mathbb{N}}$ konvergieren. Im Falle der Konvergenz gilt

$$\lim_{n \rightarrow \infty} c_n = \lim_{n \rightarrow \infty} \operatorname{Re}(c_n) + i \lim_{n \rightarrow \infty} \operatorname{Im}(c_n).$$

Beweis. Wir setzen $c_n = a_n + ib_n$, wobei $a_n, b_n \in \mathbb{R}$.

a) Die Folge $(c_n)_{n \in \mathbb{N}}$ konvergiere gegen $c = a + ib$, $a, b \in \mathbb{R}$. Dann existiert zu jedem $\varepsilon > 0$ ein $N \in \mathbb{N}$, so dass

$$|c_n - c| < \varepsilon \quad \text{für alle } n \geq N.$$

Daraus folgt für alle $n \geq N$

$$\begin{aligned} |a_n - a| &= |\operatorname{Re}(c_n - c)| \leq |c_n - c| < \varepsilon, \\ |b_n - b| &= |\operatorname{Im}(c_n - c)| \leq |c_n - c| < \varepsilon. \end{aligned}$$

Also konvergieren die beiden Folgen (a_n) und (b_n) gegen $a = \operatorname{Re}(c)$ bzw. $b = \operatorname{Im}(c)$.

b) Sei jetzt umgekehrt vorausgesetzt, dass die beiden Folgen (a_n) und (b_n) gegen a bzw. b konvergieren. Zu vorgegebenem $\varepsilon > 0$ existieren dann $N_1, N_2 \in \mathbb{N}$, so dass

$$|a_n - a| < \frac{\varepsilon}{2} \quad \text{für } n \geq N_1 \quad \text{und} \quad |b_n - b| < \frac{\varepsilon}{2} \quad \text{für } n \geq N_2.$$

Sei $c := a + ib$ und $N := \max(N_1, N_2)$. Dann gilt für alle $n \geq N$

$$|c_n - c| = |(a_n - a) + i(b_n - b)| \leq |a_n - a| + |b_n - b| < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon.$$

Also konvergiert die Folge (c_n) gegen $c = a + ib$.

Corollar. Sei $(c_n)_{n \in \mathbb{N}}$ eine konvergente Folge komplexer Zahlen. Dann konvergiert auch die konjugiert-komplexe Folge $(\bar{c}_n)_{n \in \mathbb{N}}$ und es gilt

$$\lim_{n \rightarrow \infty} \bar{c}_n = \overline{\lim_{n \rightarrow \infty} c_n}.$$

Beweis. Dies folgt daraus, dass $\operatorname{Re}(\bar{c}_n) = \operatorname{Re}(c_n)$ und $\operatorname{Im}(\bar{c}_n) = -\operatorname{Im}(c_n)$.

Definition. Eine Folge komplexer Zahlen heißt *Cauchy-Folge*, wenn zu jedem $\varepsilon > 0$ ein $N \in \mathbb{N}$ existiert, so dass

$$|c_n - c_m| < \varepsilon \quad \text{für alle } n, m \geq N.$$

Man beachte, dass diese Definition völlig mit der entsprechenden Definition für reelle Folgen aus §5 übereinstimmt.

Ähnlich wie Satz 2 beweist man

Satz 3. Eine Folge $(c_n)_{n \in \mathbb{N}}$ komplexer Zahlen ist genau dann eine Cauchy-Folge, wenn die beiden reellen Folgen $(\operatorname{Re}(c_n))_{n \in \mathbb{N}}$ und $(\operatorname{Im}(c_n))_{n \in \mathbb{N}}$ Cauchy-Folgen sind.

Da in $\mathbb{R}$ jede Cauchy-Folge konvergiert, folgt daraus

Satz 4. In $\mathbb{C}$ konvergiert jede Cauchy-Folge.

Bemerkung. Satz 4 besagt, dass $\mathbb{C}$ ein vollständiger, bewerteter Körper ist.

Satz 5. Seien $(c_n)_{n \in \mathbb{N}}$ und $(d_n)_{n \in \mathbb{N}}$ konvergente Folgen komplexer Zahlen. Dann konvergieren auch die Summenfolge $(c_n + d_n)_{n \in \mathbb{N}}$ und die Produktfolge

$(c_n d_n)_{n \in \mathbb{N}}$ und es gilt

$$\begin{aligned}\lim_{n \rightarrow \infty} (c_n + d_n) &= (\lim_{n \rightarrow \infty} c_n) + (\lim_{n \rightarrow \infty} d_n), \\ \lim_{n \rightarrow \infty} (c_n d_n) &= (\lim_{n \rightarrow \infty} c_n) (\lim_{n \rightarrow \infty} d_n).\end{aligned}$$

Ist außerdem $\lim d_n \neq 0$, so gilt $d_n \neq 0$ für $n \geq n_0$ und die Folge $(c_n/d_n)_{n \geq n_0}$ konvergiert. Für ihren Grenzwert gilt

$$\lim_{n \rightarrow \infty} \frac{c_n}{d_n} = \frac{\lim_{n \rightarrow \infty} c_n}{\lim_{n \rightarrow \infty} d_n}.$$

Der Beweis kann fast wörtlich aus §4 (Satz 3 und 4) übernommen werden.

Definition. Eine Reihe $\sum_{n=0}^{\infty} c_n$ komplexer Zahlen heißt *konvergent*, wenn die Folge der Partialsummen $s_n := \sum_{k=0}^n c_k$, $n \in \mathbb{N}$, konvergiert. Sie heißt *absolut konvergent*, wenn die Reihe $\sum_{n=0}^{\infty} |c_n|$ der Absolut-Beträge konvergiert.

Das Majoranten- und das Quotienten-Kriterium für komplexe Zahlen können genau so wie im reellen Fall (siehe §7) bewiesen werden.

Majoranten-Kriterium. Sei $\sum a_n$ eine konvergente Reihe nicht-negativer reeller Zahlen a_n . Weiter sei $(c_n)_{n \in \mathbb{N}}$ eine Folge komplexer Zahlen mit $|c_n| \leq a_n$ für alle $n \in \mathbb{N}$. Dann konvergiert die Reihe $\sum_{n=0}^{\infty} c_n$ absolut.

Quotienten-Kriterium. Sei $\sum_{n=0}^{\infty} c_n$ eine Reihe komplexer Zahlen mit $c_n \neq 0$ für $n \geq n_0$. Es gebe ein $\theta \in \mathbb{R}$ mit $0 < \theta < 1$, so dass

$$\left| \frac{c_{n+1}}{c_n} \right| \leq \theta \quad \text{für alle } n \geq n_0.$$

Dann konvergiert die Reihe $\sum_{n=0}^{\infty} c_n$ absolut.

Satz 6. Für jedes $z \in \mathbb{C}$ ist die Exponentialreihe

$$\exp(z) := \sum_{n=0}^{\infty} \frac{z^n}{n!}$$

absolut konvergent.

Beweis. Sei $z \neq 0$. Mit $c_n := z^n/n!$ gilt für alle $n \geq 2|z|$

$$\left| \frac{c_{n+1}}{c_n} \right| = \left| \frac{z^{n+1}}{(n+1)!} \cdot \frac{n!}{z^n} \right| = \frac{|z|}{n+1} \leq \frac{1}{2}.$$

Die Behauptung folgt deshalb aus dem Quotienten-Kriterium.

Wie in §8, Satz 2, zeigt man die

Abschätzung des Restglieds. Es gilt

$$\exp(z) = \sum_{n=0}^N \frac{z^n}{n!} + R_{N+1}(z),$$

wobei $|R_{N+1}(z)| \leq 2 \frac{|z|^{N+1}}{(N+1)!}$ für alle z mit $|z| \leq 1 + \frac{1}{2}N$.

Satz 7 (Funktionalgleichung der Exponentialfunktion). *Für alle $z_1, z_2 \in \mathbb{C}$ gilt*

$$\exp(z_1 + z_2) = \exp(z_1) \exp(z_2).$$

Beweis. Dies wird wie Satz 4 aus §8 bewiesen. Das ist möglich, da der dort vorangehende Satz 3 über das Cauchy-Produkt von Reihen richtig bleibt, wenn man $\sum a_n$ und $\sum b_n$ durch absolut konvergente Reihen komplexer Zahlen ersetzt. Der Beweis muss nicht abgeändert werden.

Corollar. *Für alle $z \in \mathbb{C}$ gilt $\exp(z) \neq 0$.*

Beweis. Es gilt $\exp(z)\exp(-z) = \exp(z-z) = \exp(0) = 1$. Wäre $\exp(z) = 0$, ergäbe sich daraus der Widerspruch $0 = 1$.

Bemerkung. Im Reellen hatten wir $\exp(x) > 0$ für alle $x \in \mathbb{R}$ bewiesen. Dies gilt natürlich im Komplexen nicht, da $\exp(z)$ im Allgemeinen nicht reell ist. Aber selbst wenn $\exp(z)$ reell ist, braucht es nicht positiv zu sein. So werden wir z.B. im nächsten Paragraphen beweisen, dass $\exp(i\pi) = -1$.

Satz 8. *Für jedes $z \in \mathbb{C}$ gilt $\exp(\bar{z}) = \overline{\exp(z)}$.*

Beweis. Sei $s_n(z) := \sum_{k=0}^n \frac{z^k}{k!}$ und $s_n^*(z) := \sum_{k=0}^n \frac{\bar{z}^k}{k!}$.

Nach den Rechenregeln für die Konjugation gilt für alle $n \in \mathbb{N}$

$$\overline{s_n(z)} = \overline{\sum_{k=0}^n \frac{z^k}{k!}} = \sum_{k=0}^n \overline{\left(\frac{z^k}{k!}\right)} = \sum_{k=0}^n \frac{\bar{z}^k}{k!} = s_n^*(z).$$

Aus dem Corollar zu Satz 2 folgt daher

$$\exp(\bar{z}) = \lim_{n \rightarrow \infty} s_n^*(z) = \lim_{n \rightarrow \infty} \overline{s_n(z)} = \overline{\lim_{n \rightarrow \infty} s_n(z)} = \overline{\exp(z)}, \quad \text{q.e.d.}$$

Definition. Sei D eine Teilmenge von $\mathbb{C}$. Eine Funktion $f: D \rightarrow \mathbb{C}$ heißt stetig in einem Punkt $p \in D$, falls

$$\lim_{\substack{z \rightarrow p \\ z \in D}} f(z) = f(p),$$

d.h. wenn für jede Folge $(z_n)_{n \in \mathbb{N}}$ von Punkten $z_n \in D$ mit $\lim z_n = p$ gilt $\lim_{n \rightarrow \infty} f(z_n) = f(p)$. Die Funktion f heißt stetig in D , wenn sie in jedem Punkt $p \in D$ stetig ist.

Satz 9. *Die Exponentialfunktion*

$$\exp: \mathbb{C} \longrightarrow \mathbb{C}, \quad z \mapsto \exp(z),$$

ist in ganz $\mathbb{C}$ stetig.

Beweis. Die Abschätzung des Restglieds der Exponentialreihe liefert für $N = 0$:

$$|\exp(z) - 1| \leq 2|z| \quad \text{für } |z| \leq 1.$$

Sei nun $p \in \mathbb{C}$ und (z_n) eine Folge komplexer Zahlen mit $\lim z_n = p$, also $\lim (z_n - p) = 0$. Aus der obigen Abschätzung folgt daher

$$\lim_{n \rightarrow \infty} \exp(z_n - p) = 1.$$

Mithilfe der Funktionalgleichung erhalten wir daraus

$$\lim_{n \rightarrow \infty} \exp(z_n) = \lim_{n \rightarrow \infty} \exp(p) \exp(z_n - p) = \exp(p), \quad \text{q.e.d.}$$

AUFGABEN

13.1. Sei c eine komplexe Zahl ungleich 0. Man beweise: Die Gleichung $z^2 = c$ besitzt genau zwei Lösungen. Für eine der beiden Lösungen gilt

$$\operatorname{Re}(z) = \sqrt{\frac{|c| + \operatorname{Re}(c)}{2}}, \quad \operatorname{Im}(z) = \sigma \sqrt{\frac{|c| - \operatorname{Re}(c)}{2}},$$

wobei

$$\sigma := \begin{cases} +1, & \text{falls } \operatorname{Im}(c) \geq 0, \\ -1, & \text{falls } \operatorname{Im}(c) < 0. \end{cases}$$

Die andere Lösung ist das Negative davon.

13.2. Sei $a \in \mathbb{R}$. Man zeige: Die Gleichung

$$z^2 + 2az + 1 = 0$$

hat genau dann keine reellen Lösungen, wenn $|a| < 1$. In diesem Fall hat die Gleichung zwei konjugiert komplexe Lösungen, die auf dem Einheitskreis $\{z \in \mathbb{C} : |z| = 1\}$ liegen.

13.3. Man bestimme alle komplexen Lösungen der Gleichungen

$$z^3 = 1 \quad \text{und} \quad w^6 = 1.$$

13.4. (Die elementare analytische Geometrie der Ebene sei vorausgesetzt.)
Man zeige: Für jedes $c \in \mathbb{C} \setminus \{0\}$ und jedes $\alpha \in \mathbb{R}$ ist

$$\{z \in \mathbb{C} : \operatorname{Re}(cz) = \alpha\}$$

eine Gerade in $\mathbb{C}$. Umgekehrt lässt sich jede Gerade in der komplexen Ebene $\mathbb{C}$ so darstellen.

13.5. Sei $z_1 := -1 - i$ und $z_2 := 3 + 2i$. Man bestimme eine Zahl $z_3 \in \mathbb{C}$, so dass z_1, z_2, z_3 die Ecken eines gleichseitigen Dreiecks bilden.

13.6. Man zeige:

a) Für jede Zahl $\zeta \in \mathbb{C} \setminus \{0\}$ gilt $|\bar{\zeta}/\zeta| = 1$.

b) Zu jeder Zahl $z \in \mathbb{C}$ mit $|z| = 1$ gibt es ein $\zeta \in \mathbb{C} \setminus \{0\}$ mit $z = \bar{\zeta}/\zeta$.

13.7. Man untersuche die folgenden Reihen auf Konvergenz:

$$\sum_{n=1}^{\infty} \frac{i^n}{n}, \quad \sum_{n=0}^{\infty} \left(\frac{1-i}{1+i} \right)^n, \quad \sum_{n=2}^{\infty} \frac{1}{\log(n)} \left(\frac{1-i}{1+i} \right)^n, \quad \sum_{n=2}^{\infty} n^2 \left(\frac{1-i}{2+i} \right)^n.$$

13.8. Es sei $k \geq 1$ eine natürliche Zahl und für $n \in \mathbb{N}$ seien

$$A_n \in M(k \times k, \mathbb{C}), \quad A_n = (a_{ij}^{(n)}),$$

komplexe $k \times k$ -Matrizen. Man sagt, die Folge $(A_n)_{n \in \mathbb{N}}$ konvergiere gegen die Matrix $A = (a_{ij}) \in M(k \times k, \mathbb{C})$, falls für jedes Paar $(i, j) \in \{1, \dots, k\}^2$ gilt

$$\lim_{n \rightarrow \infty} a_{ij}^{(n)} = a_{ij}.$$

Man beweise:

i) Für jede Matrix $A \in M(k \times k, \mathbb{C})$ konvergiert die Reihe

$$\exp(A) := \sum_{n=0}^{\infty} \frac{1}{n!} A^n.$$

ii) Seien $A, B \in M(k \times k, \mathbb{C})$ Matrizen mit $AB = BA$. Dann gilt

$$\exp(A+B) = \exp(A)\exp(B).$$

§ 14 Trigonometrische Funktionen

Wie bereits angekündigt, führen wir nun die trigonometrischen Funktionen mithilfe der Eulerschen Formel $e^{ix} = \cos x + i \sin x$ ein. Ihre wichtigsten Eigenschaften, wie Reihenentwicklung, Additionstheoreme und Periodizität ergeben sich daraus in einfacher Weise. Außerdem behandeln wir in diesem Paragraphen die Arcus-Funktionen, die Umkehrfunktionen der trigonometrischen Funktionen.

Definition (Cosinus, Sinus). Für $x \in \mathbb{R}$ sei

$$\cos x := \operatorname{Re}(e^{ix}),$$

$$\sin x := \operatorname{Im}(e^{ix}).$$

Es ist also $e^{ix} = \cos x + i \sin x$ (*Eulersche Formel*).

Geometrische Deutung von Cosinus und Sinus in der Gaußschen Zahlenebene. Für alle $x \in \mathbb{R}$ ist $|e^{ix}| = 1$, denn nach §13, Satz 8, gilt

$$|e^{ix}|^2 = e^{ix} \overline{e^{ix}} = e^{ix} e^{-ix} = e^0 = 1.$$

e^{ix} ist also ein Punkt des Einheitskreises der Gaußschen Ebene und $\cos x$ bzw. $\sin x$ sind die Projektionen dieses Punktes auf die reelle bzw. imaginäre Achse (Bild 14.1).

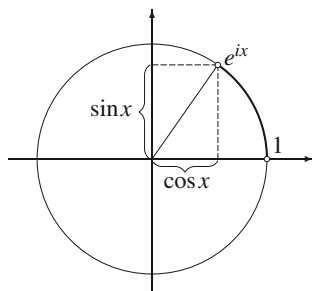


Bild 14.1

Nach Aufgabe 14.1 kann man x als orientierte Länge des Bogens von 1 nach e^{ix} mit Parameterdarstellung $t \mapsto e^{it}$, $0 \leq t \leq x$, (bzw. $0 \geq t \geq x$, falls x negativ ist) deuten.

Satz 1. Für alle $x \in \mathbb{R}$ gilt:

- a) $\cos x = \frac{1}{2}(e^{ix} + e^{-ix}), \quad \sin x = \frac{1}{2i}(e^{ix} - e^{-ix}).$
- b) $\cos(-x) = \cos x, \quad \sin(-x) = -\sin x.$
- c) $\cos^2 x + \sin^2 x = 1.$

Die Behauptungen ergeben sich unmittelbar aus der Definition.

Satz 2. Die Funktionen $\cos: \mathbb{R} \rightarrow \mathbb{R}$ und $\sin: \mathbb{R} \rightarrow \mathbb{R}$ sind auf ganz $\mathbb{R}$ stetig.

Beweis. Sei $a \in \mathbb{R}$ und (x_n) eine Folge reeller Zahlen mit $\lim x_n = a$. Daraus folgt $\lim(ix_n) = ia$, also wegen der Stetigkeit der Exponentialfunktion

$$\lim e^{ix_n} = e^{ia}.$$

Nach §13, Satz 2, gilt nun

$$\lim \cos x_n = \lim \operatorname{Re}(e^{ix_n}) = \operatorname{Re}(e^{ia}) = \cos a,$$

$$\lim \sin x_n = \lim \operatorname{Im}(e^{ix_n}) = \operatorname{Im}(e^{ia}) = \sin a.$$

Also sind $\cos$ und $\sin$ in a stetig.

Satz 3 (Additionstheoreme). Für alle $x, y \in \mathbb{R}$ gilt

$$\cos(x+y) = \cos x \cos y - \sin x \sin y,$$

$$\sin(x+y) = \sin x \cos y + \cos x \sin y.$$

Insbesondere gelten für alle $x \in \mathbb{R}$ die Verdoppelungsformeln

$$\cos(2x) = \cos^2 x - \sin^2 x = 2\cos^2 x - 1,$$

$$\sin(2x) = 2\sin x \cos x.$$

Beweis. Aus der Funktionalgleichung der Exponentialfunktion

$$e^{i(x+y)} = e^{ix+iy} = e^{ix}e^{iy}$$

ergibt sich mit der Eulerschen Formel

$$\begin{aligned} \cos(x+y) + i\sin(x+y) &= (\cos x + i\sin x)(\cos y + i\sin y) \\ &= (\cos x \cos y - \sin x \sin y) + i(\sin x \cos y + \cos x \sin y). \end{aligned}$$

Vergleicht man Real- und Imaginärteil, erhält man die Additionstheoreme für $\cos$ und $\sin$. Die Verdoppelungsformeln folgen daraus unmittelbar.

Corollar. Für alle $x, y \in \mathbb{R}$ gilt

$$\cos x - \cos y = -2 \sin \frac{x+y}{2} \sin \frac{x-y}{2},$$

$$\sin x - \sin y = 2 \cos \frac{x+y}{2} \sin \frac{x-y}{2}.$$

Beweis. Setzen wir $u := \frac{x+y}{2}$, $v := \frac{x-y}{2}$, so ist $x = u + v$ und $y = u - v$. Aus Satz 3 folgt

$$\begin{aligned} \cos x - \cos y &= \cos(u+v) - \cos(u-v) \\ &= (\cos u \cos v - \sin u \sin v) - (\cos u \cos(-v) - \sin u \sin(-v)) \\ &= -2 \sin u \sin v = -2 \sin \frac{x+y}{2} \sin \frac{x-y}{2}. \end{aligned}$$

Die zweite Gleichung ist analog zu beweisen.

Satz 4. Für alle $x \in \mathbb{R}$ gilt

$$\cos x = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k}}{(2k)!} = 1 - \frac{x^2}{2!} + \frac{x^4}{4!} \mp \dots,$$

$$\sin x = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k+1}}{(2k+1)!} = x - \frac{x^3}{3!} + \frac{x^5}{5!} \mp \dots$$

Diese Reihen konvergieren absolut für alle $x \in \mathbb{R}$.

Beweis. Die absolute Konvergenz folgt unmittelbar aus der absoluten Konvergenz der Exponentialreihe.

Für die Potenzen von i gilt

$$i^n = \begin{cases} 1, & \text{falls } n = 4m, \\ i, & \text{falls } n = 4m + 1, \\ -1, & \text{falls } n = 4m + 2, \\ -i, & \text{falls } n = 4m + 3, \end{cases} \quad (m \in \mathbb{N}).$$

Damit erhält man aus der Exponentialreihe

$$\begin{aligned} e^{ix} &= \sum_{n=0}^{\infty} \frac{(ix)^n}{n!} = \sum_{n=0}^{\infty} i^n \frac{x^n}{n!} \\ &= \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k}}{(2k)!} + i \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k+1}}{(2k+1)!}. \end{aligned}$$

Da $\cos x = \operatorname{Re}(e^{ix})$ und $\sin x = \operatorname{Im}(e^{ix})$, folgt die Behauptung.

Satz 5 (Abschätzung der Restglieder). *Es gilt*

$$\begin{aligned}\cos x &= \sum_{k=0}^n (-1)^k \frac{x^{2k}}{(2k)!} + r_{2n+2}(x), \\ \sin x &= \sum_{k=0}^n (-1)^k \frac{x^{2k+1}}{(2k+1)!} + r_{2n+3}(x),\end{aligned}$$

wobei

$$\begin{aligned}|r_{2n+2}(x)| &\leq \frac{|x|^{2n+2}}{(2n+2)!} \quad \text{für } |x| \leq 2n+3, \\ |r_{2n+3}(x)| &\leq \frac{|x|^{2n+3}}{(2n+3)!} \quad \text{für } |x| \leq 2n+4.\end{aligned}$$

Bemerkung. Die Restgliedabschätzungen sind sogar für alle $x \in \mathbb{R}$ gültig, wie später (§22) aus der Taylor-Formel folgt.

Beweis. Es ist

$$r_{2n+2}(x) = \pm \frac{x^{2n+2}}{(2n+2)!} \left(1 - \frac{x^2}{(2n+3)(2n+4)} \pm \dots \right).$$

Für $k \geq 1$ sei

$$a_k := \frac{x^{2k}}{(2n+3)(2n+4) \cdots (2n+2(k+1))}.$$

Damit ist

$$r_{2n+2}(x) = \pm \frac{x^{2n+2}}{(2n+2)!} (1 - a_1 + a_2 - a_3 \pm \dots).$$

Da

$$a_k = a_{k-1} \frac{x^2}{(2n+2k+1)(2n+2k+2)},$$

gilt für $|x| \leq 2n+3$

$$1 > a_1 > a_2 > a_3 > \dots$$

Wie beim Beweis des Leibniz'schen Konvergenzkriteriums (§7, Satz 4) folgt daraus

$$0 \leq 1 - a_1 + a_2 - a_3 \pm \dots \leq 1.$$

Deswegen ist $|r_{2n+2}(x)| \leq \frac{|x|^{2n+2}}{(2n+2)!}$.

Die Abschätzung des Restglieds von $\sin x$ ist analog zu beweisen.

Corollar. $\lim_{\substack{x \rightarrow 0 \\ x \neq 0}} \frac{\sin x}{x} = 1.$

Beweis. Wir verwenden das Restglied 3. Ordnung:

$$\sin x = x + r_3(x), \quad \text{wobei } |r_3(x)| \leq \frac{|x|^3}{3!} \text{ für } |x| \leq 4,$$

d.h.

$$|\sin x - x| \leq \frac{|x|^3}{6} \quad \text{für } |x| \leq 4.$$

Division durch x ergibt

$$\left| \frac{\sin x}{x} - 1 \right| \leq \frac{|x|^2}{6} \quad \text{für } 0 < |x| \leq 4.$$

Daraus folgt die Behauptung.

Die Zahl π

Die Zahl π wird gewöhnlich geometrisch definiert als Umfang eines Kreises vom Durchmesser 1. Eine andere, äquivalente geometrische Definition von π ist die als Fläche des Kreises mit Radius 1. Wir geben hier eine analytische Definition mithilfe der Nullstellen des Cosinus und zeigen später, dass diese Definition zu den geometrischen Definitionen äquivalent ist.

Satz 6. Die Funktion $\cos$ hat im Intervall $[0, 2]$ genau eine Nullstelle.

Zum Beweis benötigen wir drei Hilfssätze.

Hilfssatz 1. $\cos 2 \leq -\frac{1}{3}.$

Beweis. Es ist

$$\cos x = 1 - \frac{x^2}{2} + r_4(x) \quad \text{mit} \quad |r_4(x)| \leq \frac{|x|^4}{24} \quad \text{für } |x| \leq 5.$$

Speziell für $x = 2$ ergibt sich

$$\cos 2 = 1 - 2 + r_4(2) \quad \text{mit} \quad |r_4(2)| \leq \frac{16}{24} = \frac{2}{3},$$

also

$$\cos 2 \leq 1 - 2 + \frac{2}{3} = -\frac{1}{3}, \quad \text{q.e.d.}$$

Hilfssatz 2. $\sin x > 0$ für alle $x \in]0, 2]$.

Beweis. Für $x \neq 0$ können wir schreiben

$$\sin x = x + r_3(x) = x \left(1 + \frac{r_3(x)}{x} \right).$$

Nach Satz 5 ist

$$\left| \frac{r_3(x)}{x} \right| \leq \frac{|x|^2}{6} \leq \frac{4}{6} = \frac{2}{3} \quad \text{für alle } x \in]0, 2],$$

also

$$\sin x \geq x \left(1 - \frac{2}{3} \right) = \frac{x}{3} > 0 \quad \text{für alle } x \in]0, 2].$$

Hilfssatz 3. Die Funktion $\cos$ ist im Intervall $[0, 2]$ streng monoton fallend.

Beweis. Sei $0 \leq x < x' \leq 2$. Dann folgt aus Hilfssatz 2 und dem Corollar zu Satz 3

$$\cos x' - \cos x = -2 \sin \frac{x' + x}{2} \sin \frac{x' - x}{2} < 0, \quad \text{q.e.d.}$$

Beweis von Satz 6. Da $\cos 0 = 1$ und $\cos 2 \leq -\frac{1}{3}$, besitzt die Funktion $\cos$ nach dem Zwischenwertsatz im Intervall $[0, 2]$ mindestens eine Nullstelle. Nach Hilfssatz 3 gibt es nicht mehr als eine Nullstelle.

Wir können nun die Zahl π definieren.

Definition. $\frac{\pi}{2}$ ist die (eindeutig bestimmte) Nullstelle der Funktion $\cos$ im Intervall $[0, 2]$.

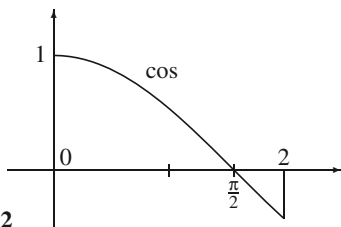


Bild 14.2

Näherungsweise Berechnung von π

Obwohl es effizientere Methoden zur Berechnung von π gibt (eine davon ist in Aufgabe 22.6 beschrieben), lässt sich die obige Definition auch direkt zur näherungsweisen Berechnung von π benutzen. Wir schreiben dazu eine ARIBAS-Funktion $\cos 20$, die den Cosinus durch den Anfang seiner Reihen-Entwicklung bis einschließlich des Gliedes der Ordnung 20 berechnet. Der Fehler ist dann nach Satz 5 kleiner als $|x|^{22}/22! < 10^{-21}|x|^{22}$.

```

function cos20(x: real): real;
var
  z, u, xx: real;
  k: integer;
begin
  z := u := 1.0;
  xx := -x*x;
  for k := 1 to 10 do
    u := u*xx/((2*k-1)*2*k);
    z := z + u;
  end;
  return z;
end.

```

Mit der Funktion `findzero` aus §11 berechnen wir nun ein Intervall der Länge 10^{-15} , in dem $\pi/2$ liegt. Dazu muss zuerst die Rechengenauigkeit auf `double_float` eingestellt werden (das entspricht in ARIBAS einer Mantissenlänge von 64 Bit; es ist $2^{-64} \approx 5.4 \cdot 10^{-20}$).

```

==> set_floatprec(double_float).
-: 64

==> eps := 10**-15.
-: 1.0000000000000000E-15

==> x0 := findzero(cos20,0,2,eps).
-: 1.57079632679489700

==> cos20(x0-eps/2).
-: 1.35479697569886180E-16

==> cos20(x0+eps/2).
-: -8.64512190806724724E-16

```

Da $1.6^{22}/22! < 3 \cdot 10^{-17}$, ist damit $\pi/2$ mit einer Genauigkeit $\pm 0.5 \cdot 10^{-15}$ ausgerechnet, und man erhält

$$\pi = 3.141592653589794 \pm 10^{-15}.$$

Satz 7 (Spezielle Werte der Exponentialfunktion).

$$e^{i\frac{\pi}{2}} = i, \quad e^{i\pi} = -1, \quad e^{i\frac{3\pi}{2}} = -i, \quad e^{2\pi i} = 1.$$

Beweis. Da $\cos \frac{\pi}{2} = 0$, ist

$$\sin^2 \frac{\pi}{2} = 1 - \cos^2 \frac{\pi}{2} = 1.$$

Nach Hilfssatz 2 ist daher $\sin \frac{\pi}{2} = +1$, also

$$e^{i\frac{\pi}{2}} = \cos \frac{\pi}{2} + i \sin \frac{\pi}{2} = i.$$

Die restlichen Behauptungen folgen wegen $e^{i\frac{n\pi}{2}} = i^n$.

Aus Satz 7 ergibt sich folgende Wertetabelle für $\sin$ und $\cos$.

x	0	$\frac{\pi}{2}$	π	$\frac{3\pi}{2}$	2π
$\sin x$	0	1	0	-1	0
$\cos x$	1	0	-1	0	1

Corollar 1. Für alle $x \in \mathbb{R}$ gilt

- a) $\cos(x + 2\pi) = \cos x$, $\sin(x + 2\pi) = \sin x$.
- b) $\cos(x + \pi) = -\cos x$, $\sin(x + \pi) = -\sin x$.
- c) $\cos x = \sin\left(\frac{\pi}{2} - x\right)$, $\sin x = \cos\left(\frac{\pi}{2} - x\right)$.

Dies folgt unmittelbar aus den Additionstheoremen und der obigen Wertetabelle.

Bemerkung. Aus dem Corollar folgt, dass man die Funktionen $\cos$ und $\sin$ nur im Intervall $[0, \frac{\pi}{4}]$ zu kennen braucht, um den Gesamtverlauf der Funktionen $\cos$ und $\sin$ zu kennen. Die Graphen von $\cos$ und $\sin$ sind in Bild 14.3 dargestellt.

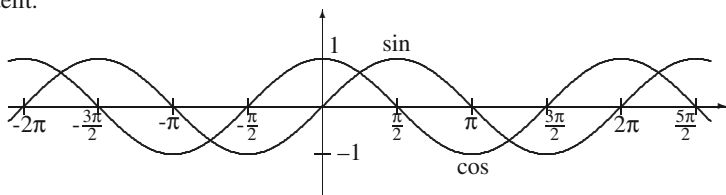


Bild 14.3 Graphen von Sinus und Cosinus

Corollar 2 (Nullstellen von Sinus und Cosinus).

- a) $\{x \in \mathbb{R} : \sin x = 0\} = \{k\pi : k \in \mathbb{Z}\}$,
- b) $\{x \in \mathbb{R} : \cos x = 0\} = \{\frac{\pi}{2} + k\pi : k \in \mathbb{Z}\}$.

Beweis. a) Nach Definition von $\frac{\pi}{2}$ und wegen $\cos(-x) = \cos x$ gilt $\cos x > 0$ für $-\frac{\pi}{2} < x < \frac{\pi}{2}$. Da $\sin x = \cos(\frac{\pi}{2} - x)$, folgt daraus

$$\sin x > 0 \quad \text{für } 0 < x < \pi.$$

Wegen $\sin(x + \pi) = -\sin x$ gilt

$$\sin x < 0 \quad \text{für } \pi < x < 2\pi.$$

Daraus folgt, dass 0 und π die einzigen Nullstellen von $\sin$ im Intervall $[0, 2\pi[$ sind. Sei nun x eine beliebige reelle Zahl mit $\sin x = 0$ und $m := \lfloor x/2\pi \rfloor$. Dann gilt

$$x = 2m\pi + \xi \quad \text{mit } 0 \leq \xi < 2\pi$$

und $\sin \xi = \sin(x - 2m\pi) = \sin x = 0$. Also ist $\xi = 0$ oder π , d.h. $x = 2m\pi$ oder $x = (2m + 1)\pi$.

Umgekehrt gilt natürlich $\sin k\pi = 0$ für alle $k \in \mathbb{Z}$.

b) Dies folgt aus a) wegen $\cos x = -\sin(x - \frac{\pi}{2})$.

Corollar 3. Für $x \in \mathbb{R}$ gilt $e^{ix} = 1$ genau dann, wenn x ein ganzzahliges Vielfaches von 2π ist.

Beweis. Wegen

$$\sin \frac{x}{2} = \frac{1}{2i} (e^{ix/2} - e^{-ix/2}) = \frac{e^{-ix/2}}{2i} (e^{ix} - 1)$$

gilt $e^{ix} = 1$ genau dann, wenn $\sin \frac{x}{2} = 0$. Die Behauptung folgt deshalb aus Corollar 2 a).

Definition (Tangens, Cotangens).

a) Die Tangensfunktion ist für $x \in \mathbb{R} \setminus \{\frac{\pi}{2} + k\pi : k \in \mathbb{Z}\}$ definiert durch

$$\tan x := \frac{\sin x}{\cos x}.$$

b) Die Cotangensfunktion ist für $x \in \mathbb{R} \setminus \{k\pi : k \in \mathbb{Z}\}$ definiert durch

$$\cot x := \frac{\cos x}{\sin x}.$$

Aus Corollar 1 c) folgt $\cot x = \tan(\frac{\pi}{2} - x)$.

Bemerkung. In der älteren Literatur werden noch die Funktionen Secans und Cosecans definiert durch

$$\sec x := \frac{1}{\cos x}, \quad \operatorname{cosec} x := \frac{1}{\sin x}.$$

Der Graph des Tangens ist in Bild 14.4 dargestellt.

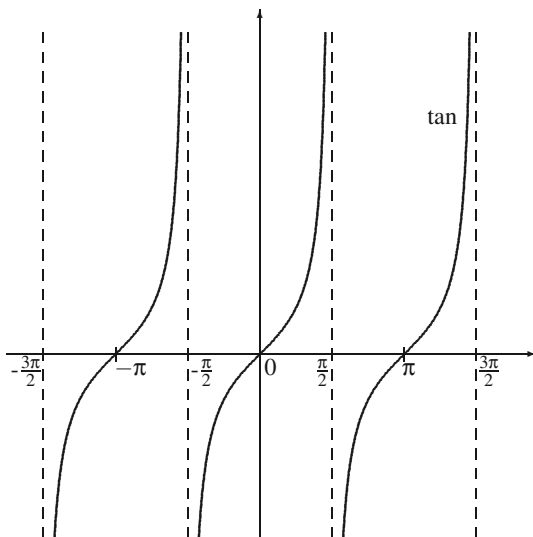


Bild 14.4 Tangens

Umkehrfunktionen der trigonometrischen Funktionen

Satz 8 und Definition.

a) Die Funktion $\cos$ ist im Intervall $[0, \pi]$ streng monoton fallend und bildet dieses Intervall bijektiv auf $[-1, 1]$ ab. Die Umkehrfunktion

$$\arccos: [-1, 1] \rightarrow \mathbb{R}$$

heißt Arcus-Cosinus.

b) Die Funktion $\sin$ ist im Intervall $[-\frac{\pi}{2}, \frac{\pi}{2}]$ streng monoton wachsend und bildet dieses Intervall bijektiv auf $[-1, 1]$ ab. Die Umkehrfunktion

$$\arcsin: [-1, 1] \rightarrow \mathbb{R}$$

heißt Arcus-Sinus.

c) Die Funktion $\tan$ ist im Intervall $]-\frac{\pi}{2}, \frac{\pi}{2}[$ streng monoton wachsend und bildet dieses Intervall bijektiv auf $\mathbb{R}$ ab. Die Umkehrfunktion

$$\arctan: \mathbb{R} \rightarrow \mathbb{R}$$

heißt Arcus-Tangens.

Beweis

a) Nach Hilfssatz 3 ist $\cos$ in $[0, 2\pi]$, insbesondere in $[0, \frac{\pi}{2}]$ streng monoton fallend. Da $\cos x = -\cos(\pi - x)$, ist $\cos$ auch in $[\frac{\pi}{2}, \pi]$ streng monoton fallend. Nach §12, Satz 1, bildet daher $\cos$ das Intervall $[0, \pi]$ bijektiv auf $[\cos \pi, \cos 0] = [-1, 1]$ ab.

b) Da $\sin x = \cos(\frac{\pi}{2} - x)$, folgt aus a), dass $\sin$ im Intervall $[-\frac{\pi}{2}, \frac{\pi}{2}]$ streng monoton wächst und daher dieses Intervall bijektiv auf $[\sin(-\frac{\pi}{2}), \sin(\frac{\pi}{2})] = [-1, 1]$ abbildet.

c) i) Sei $0 \leq x < x' < \frac{\pi}{2}$. Dann gilt $\sin x < \sin x'$ und $\cos x > \cos x' > 0$. Daraus folgt

$$\tan x = \frac{\sin x}{\cos x} < \frac{\sin x'}{\cos x'} = \tan x',$$

$\tan$ ist also in $[0, \frac{\pi}{2}[$ streng monoton wachsend. Weil $\tan(-x) = -\tan x$, wächst $\tan$ auch in $]-\frac{\pi}{2}, 0]$, d.h. im ganzen Intervall $]-\frac{\pi}{2}, \frac{\pi}{2}[$ streng monoton.

ii) Wir zeigen jetzt, dass $\lim_{x \nearrow \frac{\pi}{2}} \tan x = \infty$.

Sei $(x_n)_{n \in \mathbb{N}}$ eine Folge mit $x_n < \frac{\pi}{2}$ und $\lim x_n = \frac{\pi}{2}$. Wir dürfen annehmen, dass $x_n > 0$ für alle n . Dann ist auch

$$y_n := \frac{\cos x_n}{\sin x_n} > 0 \quad \text{für alle } n \in \mathbb{N}$$

und

$$\lim y_n = \frac{\lim \cos x_n}{\lim \sin x_n} = \frac{\cos \frac{\pi}{2}}{\sin \frac{\pi}{2}} = \frac{0}{1} = 0.$$

Daraus folgt (§4, Satz 9)

$$\lim \tan x_n = \lim \frac{1}{y_n} = \infty, \quad \text{q.e.d.}$$

iii) Wegen $\tan(-x) = -\tan x$ folgt aus ii)

$$\lim_{x \searrow -\frac{\pi}{2}} \tan x = -\infty.$$

iv) Mithilfe von §12, Satz 1, ergibt sich aus i)–iii), dass $\tan$ das Intervall $]-\frac{\pi}{2}, \frac{\pi}{2}[$ bijektiv auf $\mathbb{R}$ abbildet.

Die Graphen der Arcus-Funktionen sind in den Bildern 14.5–14.7 dargestellt.

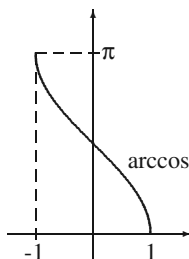


Bild 14.5 Arcus cosinus

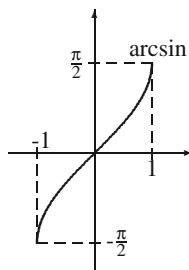


Bild 14.6 Arcus sinus

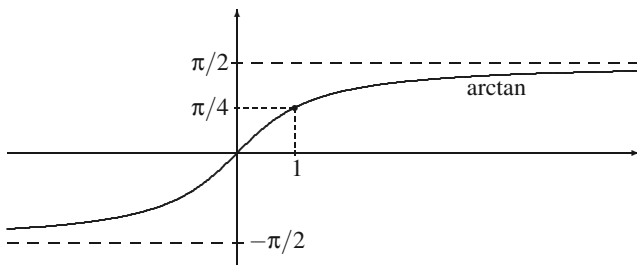


Bild 14.7 Arcus tangens

Bemerkung. Die in Satz 8 definierten Funktionen nennt man auch die *Hauptzweige* von $\arccos$, $\arcsin$ und $\arctan$. Für beliebiges $k \in \mathbb{Z}$ gilt:

- a) $\cos$ bildet $[k\pi, (k+1)\pi]$ bijektiv auf $[-1, 1]$ ab,
- b) $\sin$ bildet $[-\frac{\pi}{2} + k\pi, \frac{\pi}{2} + k\pi]$ bijektiv auf $[-1, 1]$ ab,
- c) $\tan$ bildet $]-\frac{\pi}{2} + k\pi, \frac{\pi}{2} + k\pi[$ bijektiv auf $\mathbb{R}$ ab.

Die zugehörigen Umkehrfunktionen

$$\arccos_k: [-1, 1] \rightarrow \mathbb{R},$$

$$\arcsin_k: [-1, 1] \rightarrow \mathbb{R},$$

$$\arctan_k: \mathbb{R} \rightarrow \mathbb{R}$$

heißten für $k \neq 0$ *Nebenzweige* von $\arccos$, $\arcsin$ bzw. $\arctan$.

Polarkoordinaten

Jede komplexe Zahl z wird durch zwei reelle Zahlen, den Realteil und den Imaginärteil von z dargestellt. Dies sind die kartesischen Koordinaten von z in der Gauß'schen Zahlenebene. Eine andere oft nützliche Darstellung wird durch die Polarkoordinaten gegeben.

Satz 9 (Polarkoordinaten). *Jede komplexe Zahl z lässt sich schreiben als*

$$z = r \cdot e^{i\varphi},$$

wobei $\varphi \in \mathbb{R}$ und $r = |z| \in \mathbb{R}_+$. Für $z \neq 0$ ist φ bis auf ein ganzzahliges Vielfaches von 2π eindeutig bestimmt.

Bemerkung. Die Zahl φ ist der Winkel (im Bogenmaß) zwischen der positiven reellen Achse und dem Ortsvektor von z (Bild 14.8). Man nennt φ auch das *Argument* der komplexen Zahl $z = r \cdot e^{i\varphi}$.

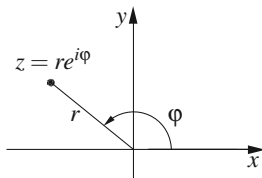


Bild 14.8

Beweis. Für $z = 0$ ist $z = 0 \cdot e^{i\varphi}$ mit beliebigem φ . Sei jetzt $z \neq 0$, $r := |z|$ und $\zeta := \frac{z}{r}$. Dann ist $|\zeta| = 1$. Sind ξ und η Real- und Imaginärteil von ζ , d.h. $\zeta = \xi + i\eta$, so gilt also $\xi^2 + \eta^2 = 1$ und $|\xi| \leq 1$. Deshalb ist

$$\alpha := \arccos \xi$$

definiert. Da $\cos \alpha = \xi$, folgt

$$\sin \alpha = \pm \sqrt{1 - \xi^2} = \pm \eta.$$

Wir setzen $\varphi := \alpha$, falls $\sin \alpha = \eta$ und $\varphi := -\alpha$, falls $\sin \alpha = -\eta$. In jedem Fall ist dann

$$e^{i\varphi} = \cos \varphi + i \sin \varphi = \xi + i\eta = \zeta.$$

Damit gilt $z = re^{i\varphi}$. Die Eindeutigkeit von φ bis auf ein Vielfaches von 2π folgt aus Corollar 3 zu Satz 7. Denn $e^{i\varphi} = e^{i\psi} = \zeta$ impliziert $e^{i(\varphi-\psi)} = 1$, also $\varphi - \psi = 2k\pi$ mit einer ganzen Zahl k .

Bemerkung. Satz 9 erlaubt eine einfache Interpretation der Multiplikation komplexer Zahlen. Sei $z = r_1 e^{i\varphi}$ und $w = r_2 e^{i\psi}$. Dann ist $zw = r_1 r_2 e^{i(\varphi+\psi)}$. Man erhält also das Produkt zweier komplexer Zahlen, indem man ihre Beträge multipliziert und ihre Argumente addiert (Bild 14.9).

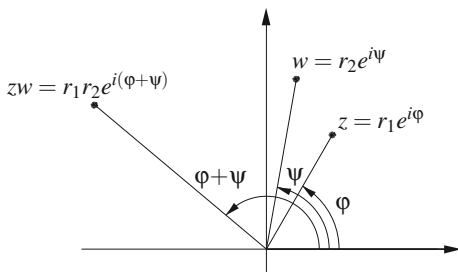


Bild 14.9 Zur Multiplikation komplexer Zahlen

Corollar (*n*-te Einheitswurzeln). Sei n eine natürliche Zahl ≥ 2 . Die Gleichung $z^n = 1$ hat genau n komplexe Lösungen, nämlich $z = w_k$, wobei

$$w_k = e^{i\frac{2k\pi}{n}}, \quad k = 0, 1, \dots, n-1.$$

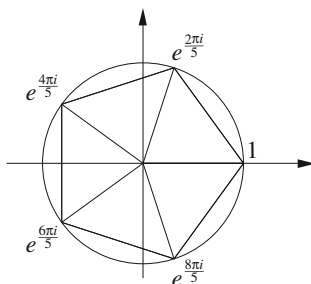
Beweis des Corollars. Die Zahl $z \in \mathbb{C}$ genüge der Gleichung $z^n = 1$. Wir können z darstellen als $z = re^{i\varphi}$ mit $0 \leq \varphi < 2\pi$ und $r \geq 0$. Da

$$1 = |z^n| = |z|^n = r^n,$$

ist $r = 1$, also

$$z^n = (e^{i\varphi})^n = e^{in\varphi} = 1.$$

Nach Corollar 3 zu Satz 7 existiert ein $k \in \mathbb{Z}$ mit $n\varphi = 2k\pi$, d.h. $\varphi = \frac{2k\pi}{n}$. Wegen $0 \leq \varphi < 2\pi$ ist $0 \leq k < n$ und $z = w_k$.

**Bild 14.10** Fünfte Einheitswurzeln

Umgekehrt gilt für jedes k

$$w_k^n = \left(e^{i\frac{2k\pi}{n}} \right)^n = e^{i2k\pi} = 1.$$

Anwendung

(14.1) Reguläres n -Eck. Die n -ten Einheitswurzeln bilden die Ecken eines dem Einheitskreis einbeschriebenen gleichseitigen n -Ecks (s. Bild 14.10). Die k -te Seite ist die Strecke von $e^{2\pi i(k-1)/n}$ nach $e^{2\pi i k/n}$ ($k = 1, \dots, n$) und hat die Länge

$$\begin{aligned} s_n &= \left| e^{2\pi i k/n} - e^{2\pi i(k-1)/n} \right| = \left| e^{2\pi i(k-1/2)/n} (e^{\pi i/n} - e^{-\pi i/n}) \right| \\ &= \left| e^{\pi i/n} - e^{-\pi i/n} \right| = 2 \sin(\pi/n). \end{aligned}$$

Infolgedessen ist der Umfang des regulären n -Ecks gleich

$$L_n = 2n \sin(\pi/n).$$

Lässt man n gegen ∞ streben, so schmiegen sich die regulären n -Ecke immer mehr dem Einheitskreis an. Der Grenzwert der Längen L_n ist als Umfang des Einheitskreises definiert. Es gilt

$$\begin{aligned} \lim_{n \rightarrow \infty} L_n &= \lim_{n \rightarrow \infty} 2n \sin(\pi/n) = \lim_{n \rightarrow \infty} 2\pi \frac{\sin(\pi/n)}{\pi/n} \\ &= 2\pi \lim_{x \rightarrow 0} \frac{\sin x}{x} = 2\pi \end{aligned}$$

nach dem Corollar zu Satz 5. Der Umfang des Einheitskreises ist also gleich 2π . Damit haben wir den Anschluss der analytischen Definition von π an die geometrische Definition hergestellt.

AUFGABEN

14.1. Sei x eine reelle Zahl und n eine natürliche Zahl ≥ 1 . Die Punkte $A_k^{(n)}$ auf dem Einheitskreis der komplexen Ebene seien wie folgt definiert:

$$A_k^{(n)} := e^{i\frac{k}{n}x}, \quad k = 0, 1, \dots, n.$$

Sei L_n die Länge des Polygonzugs $A_0^{(n)} A_1^{(n)} \dots A_n^{(n)}$, d.h.

$$L_n = \sum_{k=1}^n \left| A_k^{(n)} - A_{k-1}^{(n)} \right|.$$

Man beweise:

- a) $L_n = 2n \left| \sin \frac{x}{2n} \right|,$
 b) $\lim_{n \rightarrow \infty} 2n \sin \frac{x}{2n} = x.$

14.2. Es sei U_n der Umfang des dem Einheitskreis umschriebenen regulären n -Ecks. Man beweise:

$$U_n = 2n \tan \frac{\pi}{n}.$$

14.3. Man beweise für alle $x, y \in \mathbb{R}$, für die $\tan x$, $\tan y$ und $\tan(x+y)$ definiert sind, das Additionstheorem des Tangens

$$\tan(x+y) = \frac{\tan x + \tan y}{1 - \tan x \tan y}.$$

14.4. Man beweise folgende Halbierungs-Formeln für die trigonometrischen Funktionen:

a) $\cos \frac{\alpha}{2} = \sqrt{\frac{1 + \cos \alpha}{2}} \quad \text{für } |\alpha| \leq \pi,$

b) $\sin \frac{\alpha}{2} = \frac{\sin \alpha}{2} \sqrt{\frac{2}{1 + \sqrt{1 - \sin^2 \alpha}}} \quad \text{für } |\alpha| \leq \pi/2,$

c) $\tan \frac{\alpha}{2} = \frac{\tan \alpha}{1 + \sqrt{1 + \tan^2 \alpha}} \quad \text{für } |\alpha| < \pi/2.$

14.5. Sei $x \in \mathbb{R}$. Die Folge $(x_n)_{n \in \mathbb{N}}$ werde rekursiv wie folgt definiert:

$$x_0 := x, \quad x_{n+1} := \frac{x_n}{1 + \sqrt{1 + x_n^2}}.$$

Man zeige:

$$\lim_{n \rightarrow \infty} (2^n x_n) = \arctan x.$$

Bemerkung. Vergleiche dazu Aufgabe 12.5, wo eine ähnliche Folge für den Logarithmus gegeben wird.

14.6. Man beweise die Korrektheit der folgenden Wertetabelle für $\sin x$, $\cos x$, $\tan x$.

x	$\frac{\pi}{3}$	$\frac{\pi}{4}$	$\frac{\pi}{5}$	$\frac{\pi}{6}$
$\sin x$	$\frac{1}{2}\sqrt{3}$	$\frac{1}{2}\sqrt{2}$	$\frac{1}{4}\sqrt{10 - 2\sqrt{5}}$	$\frac{1}{2}$
$\cos x$	$\frac{1}{2}$	$\frac{1}{2}\sqrt{2}$	$\frac{1}{4}(1 + \sqrt{5})$	$\frac{1}{2}\sqrt{3}$
$\tan x$	$\sqrt{3}$	1	$\sqrt{5 - 2\sqrt{5}}$	$\frac{1}{3}\sqrt{3}$

14.7. Sei x eine reelle Zahl. Man beweise

$$\frac{1 + ix}{1 - ix} = e^{2i\varphi},$$

wobei $\varphi = \arctan x$.

14.8. Man beweise für alle $x \in \mathbb{R} \setminus \{n\pi/2 : n \in \mathbb{Z}\}$

$$\tan x = \cot x - 2 \cot 2x.$$

14.9. Für $-1 \leq x \leq 1$ und $n \in \mathbb{N}$ sei

$$T_n(x) := \cos(n \arccos x).$$

Man zeige:

- T_n ist ein Polynom n -ten Grades in x mit ganzzahligen Koeffizienten. (T_n heißt n -tes Tschebyscheff-Polynom.)
- Es gilt die Rekursionsformel $T_{n+1}(x) = 2xT_n(x) - T_{n-1}(x)$.

14.10. Sei x eine reelle Zahl, $x \neq (2k+1)\pi$ für alle $k \in \mathbb{Z}$.

Man beweise: Ist $u := \tan \frac{x}{2}$, so gilt

$$\sin x = \frac{2u}{1+u^2}, \quad \cos x = \frac{1-u^2}{1+u^2}.$$

14.11. Sei $n \geq 2$. Man beweise die Identität

$$2^{n-1} \prod_{k=1}^{n-1} \sin \frac{k\pi}{n} = n.$$

Anleitung. Die Behauptung ist äquivalent zu $\prod_{k=1}^{n-1} (1 - e^{2\pi i k/n}) = n$. Man verwende die Polynomgleichung

$$X^n - 1 = \prod_{k=0}^{n-1} (X - \zeta_n^k), \quad \zeta_n := e^{2\pi i/n}.$$

14.12. Die Funktionen Cosinus und Sinus werden im Komplexen wie folgt definiert: Für $z \in \mathbb{C}$ sei

$$\cos z := \frac{1}{2}(e^{iz} + e^{-iz}), \quad \sin z := \frac{1}{2i}(e^{iz} - e^{-iz}).$$

Man zeige für alle $x, y \in \mathbb{R}$

$$\cos(x+iy) = \cos x \cosh y - i \sin x \sinh y$$

$$\sin(x+iy) = \sin x \cosh y + i \cos x \sinh y$$

(Die Funktionen $\cosh$ und $\sinh$ wurden in Aufgabe 10.1 definiert.)

14.13. Seien z_1 und z_2 zwei komplexe Zahlen mit $\sin z_1 = \sin z_2$.

Man zeige: Es gibt eine ganze Zahl $n \in \mathbb{Z}$, so dass

$$z_1 = z_2 + 2n\pi \quad \text{oder} \quad z_1 = -z_2 + (2n+1)\pi.$$

§ 15 Differentiation

Wir definieren jetzt den Differentialquotienten (oder die Ableitung) einer Funktion als Limes der Differenzenquotienten und beweisen die wichtigsten Rechenregeln für die Ableitung, wie Produkt-, Quotienten- und Ketten-Regel sowie die Formel für die Ableitung der Umkehrfunktion. Damit ist es dann ein leichtes, die Ableitungen aller bisher besprochenen Funktionen zu berechnen.

Definition. Sei $V \subset \mathbb{R}$ und $f: V \longrightarrow \mathbb{R}$ eine Funktion. f heißt in einem Punkt $x \in V$ *differenzierbar*, falls der Grenzwert

$$f'(x) := \lim_{\substack{\xi \rightarrow x \\ \xi \in V \setminus \{x\}}} \frac{f(\xi) - f(x)}{\xi - x}$$

existiert. (Insbesondere wird vorausgesetzt, dass es mindestens eine Folge $\xi_n \in V \setminus \{x\}$ mit $\lim_{n \rightarrow \infty} \xi_n = x$ gibt, d.h. dass x ein Häufungspunkt von V ist, vgl. die Definition in § 9, Seite 94. Dies ist z.B. stets der Fall, wenn V ein Intervall ist, das aus mehr als einem Punkt besteht.)

Der Grenzwert $f'(x)$ heißt *Differentialquotient* oder *Ableitung* von f im Punkte x . Die Funktion f heißt *differenzierbar* in V , falls f in jedem Punkt $x \in V$ differenzierbar ist.

Bemerkung. Man kann den Differentialquotienten auch darstellen als

$$f'(x) = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}.$$

Dabei sind natürlich bei der Limesbildung nur solche Folgen (h_n) mit $\lim h_n = 0$ zugelassen, für die $h_n \neq 0$ und $x + h_n \in V$ für alle n .

Geometrische Interpretation des Differentialquotienten

Der *Differenzenquotient* $\frac{f(\xi) - f(x)}{\xi - x}$ ist die Steigung der Sekante des Graphen von f durch die Punkte $(x, f(x))$ und $(\xi, f(\xi))$, vgl. Bild 15.1. Beim Grenzübergang $\xi \rightarrow x$ geht die Sekante in die Tangente an den Graphen von f im Punkt $(x, f(x))$ über. $f'(x)$ ist also (im Falle der Existenz) die Steigung der Tangente im Punkt $(x, f(x))$.

Bezeichnung. Man schreibt auch $\frac{df(x)}{dx}$ für $f'(x)$.

Diese Schreibweise erklärt sich so: Schreibt man $\frac{\Delta f(x)}{\Delta x}$ für den Differenzen-

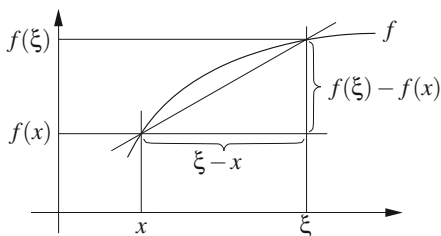


Bild 15.1

quotienten $\frac{f(\xi) - f(x)}{\xi - x}$, so wird damit

$$\frac{df(x)}{dx} = \lim_{\Delta x \rightarrow 0} \frac{\Delta f(x)}{\Delta x}.$$

Im Gegensatz zum Differenzenquotienten ist jedoch der Differentialquotient $\frac{df(x)}{dx}$ nicht der Quotient zweier reeller Zahlen $df(x)$ und dx . Die Schreibweise $\frac{df(x)}{dx}$ ist auch insofern problematisch, dass der Buchstabe x in Zähler und Nenner eine verschiedene Bedeutung hat. Ist z.B. $x = 0$, so kann man zwar $f'(0)$, aber nicht $\frac{df(0)}{d0}$ schreiben. In diesem Fall verwendet man die Schreibweisen

$$\frac{df}{dx}(0) \quad \text{oder} \quad \left. \frac{df(x)}{dx} \right|_{x=0}.$$

Beispiele

(15.1) Für eine konstante Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = c$, gilt

$$f'(x) = \lim_{\substack{\xi \rightarrow x \\ \xi \neq x}} \frac{f(\xi) - f(x)}{\xi - x} = \lim_{\substack{\xi \rightarrow x \\ \xi \neq x}} \frac{c - c}{\xi - x} = 0.$$

(15.2) $f: \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = cx$, ($c \in \mathbb{R}$).

$$f'(x) = \lim_{\substack{\xi \rightarrow x \\ \xi \neq x}} \frac{f(\xi) - f(x)}{\xi - x} = \lim_{\substack{\xi \rightarrow x \\ \xi \neq x}} \frac{c\xi - cx}{\xi - x} = c.$$

(15.3) $f: \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = x^2$.

$$\begin{aligned} f'(x) &= \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} = \lim_{h \rightarrow 0} \frac{(x+h)^2 - x^2}{h} \\ &= \lim_{h \rightarrow 0} \frac{2xh + h^2}{h} = \lim_{h \rightarrow 0} (2x + h) = 2x. \end{aligned}$$

$$(15.4) \quad f: \mathbb{R}^* \rightarrow \mathbb{R}, f(x) = \frac{1}{x}.$$

$$\begin{aligned} f'(x) &= \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} = \lim_{h \rightarrow 0} \frac{1}{h} \left(\frac{1}{x+h} - \frac{1}{x} \right) \\ &= \lim_{h \rightarrow 0} \frac{x - (x+h)}{h(x+h)x} = \lim_{h \rightarrow 0} \frac{-1}{(x+h)x} = -\frac{1}{x^2}, \end{aligned}$$

also

$$\frac{d}{dx} \left(\frac{1}{x} \right) = -\frac{1}{x^2}.$$

$$(15.5) \quad \exp: \mathbb{R} \rightarrow \mathbb{R}.$$

Unter Benutzung von Beispiel (12.7) erhält man

$$\begin{aligned} \exp'(x) &= \lim_{h \rightarrow 0} \frac{\exp(x+h) - \exp(x)}{h} = \lim_{h \rightarrow 0} \exp(x) \frac{\exp(h) - 1}{h} \\ &= \exp(x) \lim_{h \rightarrow 0} \frac{\exp(h) - 1}{h} = \exp(x). \end{aligned}$$

Die Exponentialfunktion besitzt also die merkwürdige Eigenschaft, sich bei Differentiation zu reproduzieren.

$$(15.6) \quad \sin: \mathbb{R} \rightarrow \mathbb{R}.$$

Mithilfe von §14, Corollar zu Satz 3 erhalten wir

$$\begin{aligned} \sin'(x) &= \lim_{h \rightarrow 0} \frac{\sin(x+h) - \sin x}{h} = \lim_{h \rightarrow 0} \frac{2 \cos(x + \frac{h}{2}) \sin \frac{h}{2}}{h} \\ &= \left(\lim_{h \rightarrow 0} \cos(x + \frac{h}{2}) \right) \left(\lim_{h \rightarrow 0} \frac{\sin \frac{h}{2}}{\frac{h}{2}} \right). \end{aligned}$$

Da $\cos$ stetig ist, gilt $\lim_{h \rightarrow 0} \cos(x + \frac{h}{2}) = \cos x$ und nach §14, Corollar zu Satz 5,

ist $\lim_{h \rightarrow 0} \frac{\sin \frac{h}{2}}{\frac{h}{2}} = 1$. Damit folgt

$$\sin'(x) = \cos x.$$

$$(15.7) \quad \cos: \mathbb{R} \rightarrow \mathbb{R}.$$

Analog zum vorigen Beispiel schließt man

$$\cos'(x) = \lim_{h \rightarrow 0} \frac{\cos(x+h) - \cos(x)}{h} = \lim_{h \rightarrow 0} \frac{-2 \sin(x + \frac{h}{2}) \sin \frac{h}{2}}{h}$$

$$= - \left(\lim_{h \rightarrow 0} \sin\left(x + \frac{h}{2}\right) \right) \left(\lim_{h \rightarrow 0} \frac{\sin \frac{h}{2}}{\frac{h}{2}} \right) = -\sin x,$$

also

$$\cos'(x) = -\sin(x).$$

(15.8) Die Ableitung komplexwertiger Funktionen wird ebenso definiert wie für reellwertige Funktionen. Als Beispiel betrachten wir die Funktion

$$f: \mathbb{R} \longrightarrow \mathbb{C}, \quad x \mapsto f(x) := e^{\lambda x},$$

wobei $\lambda \in \mathbb{C}$ eine komplexe Konstante ist (die Variable x ist jedoch reell).

Behauptung. $\frac{de^{\lambda x}}{dx} = \lambda e^{\lambda x}.$

Beweis. Wir verwenden die Restgliedabschätzung der Exponentialreihe im Komplexen (§13, Seite 142)

$$|\exp(\lambda x) - (1 + \lambda x)| \leq |\lambda x|^2 \quad \text{für } |\lambda x| \leq 3/2.$$

Division durch $x \neq 0$ ergibt

$$\left| \frac{e^{\lambda x} - 1}{x} - \lambda \right| \leq |\lambda^2 x| \quad \text{für } |\lambda x| \leq 3/2,$$

also $\lim_{x \rightarrow 0} \frac{e^{\lambda x} - 1}{x} = \lambda$. Daraus folgt

$$\frac{de^{\lambda x}}{dx} = \lim_{h \rightarrow 0} \frac{1}{h} (e^{\lambda(x+h)} - e^{\lambda x}) = e^{\lambda x} \lim_{h \rightarrow 0} \frac{e^{\lambda h} - 1}{h} = \lambda e^{\lambda x}, \quad \text{q.e.d.}$$

Folgerung. Speziell für $\lambda = i$ hat man $\frac{d}{dx} e^{ix} = i e^{ix}$. Setzt man darin die Eulersche Formel ein, so folgt

$$\frac{d}{dx} (\cos x + i \sin x) = i (\cos x + i \sin x) = -\sin x + i \cos x.$$

Durch Vergleich der Real- und Imaginärteile erhält man einen neuen Beweis der bereits in (15.6) und (15.7) bewiesenen Formeln

$$\cos'(x) = -\sin x, \quad \sin'(x) = \cos x.$$

(15.9) Wir betrachten die Funktion $\text{abs}: \mathbb{R} \rightarrow \mathbb{R}$ (vgl. Bild 10.1).

Behauptung. $\text{abs}'(0)$ existiert nicht.

Beweis. Sei $h_n = (-1)^n \frac{1}{n}$, ($n \geq 1$). Es gilt $\lim h_n = 0$.

$$q_n := \frac{\text{abs}(0 + h_n) - \text{abs}(0)}{h_n} = \frac{\frac{1}{n} - 0}{(-1)^n \frac{1}{n}} = (-1)^n.$$

$\lim_{n \rightarrow \infty} q_n$ existiert nicht, also ist die Funktion abs im Nullpunkt nicht differenzierbar.

Bemerkung. Sei $x \in V \subset \mathbb{R}$ und $f: V \rightarrow \mathbb{R}$ eine Funktion. f heißt im Punkt x von rechts differenzierbar, falls der Grenzwert

$$f'_+(x) := \lim_{\xi \searrow x} \frac{f(\xi) - f(x)}{\xi - x}$$

existiert. Die Funktion f heißt in x von links differenzierbar, falls

$$f'_-(x) := \lim_{\xi \nearrow x} \frac{f(\xi) - f(x)}{\xi - x}$$

existiert.

Die Funktion abs ist im Nullpunkt von rechts und von links differenzierbar, und zwar gilt $\text{abs}'_+(0) = +1$, $\text{abs}'_-(0) = -1$.

Satz 1 (Lineare Approximierbarkeit). Sei $V \subset \mathbb{R}$ und $a \in V$ ein Häufungspunkt von V . Eine Funktion $f: V \rightarrow \mathbb{R}$ ist genau dann im Punkt a differenzierbar, wenn es eine Konstante $c \in \mathbb{R}$ gibt, so dass

$$f(x) = f(a) + c(x - a) + \varphi(x), \quad (x \in V),$$

wobei φ eine Funktion ist, für die gilt

$$\lim_{\substack{x \rightarrow a \\ x \neq a}} \frac{\varphi(x)}{x - a} = 0.$$

In diesem Fall ist $c = f'(a)$.

Bemerkung. Der Satz drückt aus, dass die Differenzierbarkeit von f im Punkt a gleichbedeutend mit der Approximierbarkeit durch eine affin-lineare Funktion ist. Mit den obigen Bezeichnungen ist diese affin-lineare Funktion

$$L(x) = f(a) + c(x - a).$$

Der Graph von L ist die Tangente an den Graphen von f im Punkt $(a, f(a))$, siehe Bild 15.2. Unter Benutzung des Landauschen o -Symbols (definiert in § 12) lässt sich schreiben

$$f(x) = f(a) + c(x - a) + o(|x - a|) \quad \text{für } x \rightarrow a.$$

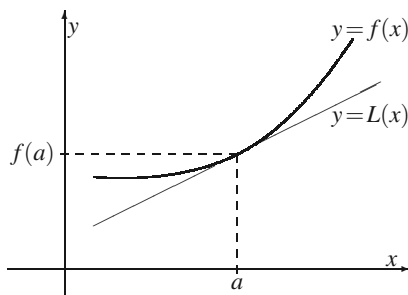


Bild 15.2 Affin-lineare Approximation

Beweis.

a) Sei zunächst vorausgesetzt, dass f in a differenzierbar ist und $c := f'(a)$. Wir definieren die Funktion φ durch

$$f(x) = f(a) + c(x-a) + \varphi(x).$$

Dann gilt

$$\frac{\varphi(x)}{x-a} = \frac{f(x) - f(a)}{x-a} - f'(a),$$

also $\lim_{x \rightarrow a} \frac{\varphi(x)}{x-a} = 0$.

b) Es sei nun umgekehrt vorausgesetzt, dass für f die Darstellung

$$f(x) = f(a) + c(x-a) + \varphi(x)$$

mit $\lim_{x \rightarrow a} \frac{\varphi(x)}{x-a} = 0$ besteht. Dann ist

$$\lim_{x \rightarrow a} \left(\frac{f(x) - f(a)}{x-a} - c \right) = \lim_{x \rightarrow a} \frac{\varphi(x)}{x-a} = 0,$$

also

$$\lim_{x \rightarrow a} \frac{f(x) - f(a)}{x-a} = c,$$

d.h. f ist in a differenzierbar und $f'(a) = c$.

Corollar. Ist die Funktion $f: V \rightarrow \mathbb{R}$ im Punkt $a \in V$ differenzierbar, so ist sie in a auch stetig.

Beweis. Wir benutzen die Darstellung von f aus Satz 1. Es gilt $\lim_{x \rightarrow a} \varphi(x) = 0$,

also

$$\lim_{x \rightarrow a} f(x) = f(a) + \lim_{x \rightarrow a} (c(x-a) + \varphi(x)) = f(a), \quad \text{q.e.d.}$$

Differentiations-Regeln

In den meisten Fällen verwendet man bei der Berechnung der Ableitung einer Funktion nicht direkt die Limes-Definition, sondern führt die Ableitung mit gewissen Regeln auf schon bekannte Fälle zurück. Diese Regeln, wie Produkt- und Quotientenregel, Kettenregel und den Satz über die Ableitung der Umkehrfunktion werden wir jetzt beweisen und Beispiele besprechen.

Satz 2 (Algebraische Operationen). *Seien $f, g: V \rightarrow \mathbb{R}$ in $x \in V$ differenzierbare Funktionen und $\lambda \in \mathbb{R}$. Dann sind auch die Funktionen*

$$f + g, \lambda f, fg: V \rightarrow \mathbb{R}$$

in x differenzierbar und es gelten die Rechenregeln:

a) Linearität

$$\begin{aligned}(f + g)'(x) &= f'(x) + g'(x), \\ (\lambda f)'(x) &= \lambda f'(x).\end{aligned}$$

b) Produktregel

$$(fg)'(x) = f'(x)g(x) + f(x)g'(x).$$

c) Quotientenregel

Ist $g(\xi) \neq 0$ für alle $\xi \in V$, so ist auch die Funktion $(f/g): V \rightarrow \mathbb{R}$ in x differenzierbar mit

$$\left(\frac{f}{g}\right)'(x) = \frac{f'(x)g(x) - f(x)g'(x)}{g(x)^2}.$$

Beweis.

a) Dies folgt unmittelbar aus den Rechenregeln für Grenzwerte von Folgen.

b) *Produktregel.*

$$\begin{aligned}(fg)'(x) &= \lim_{h \rightarrow 0} \frac{f(x+h)g(x+h) - f(x)g(x)}{h} \\ &= \lim_{h \rightarrow 0} \frac{1}{h} \{f(x+h)(g(x+h) - g(x)) + (f(x+h) - f(x))g(x)\}\end{aligned}$$

$$\begin{aligned}
&= \lim_{h \rightarrow 0} f(x+h) \frac{g(x+h) - g(x)}{h} + \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} g(x) \\
&= f(x)g'(x) + f'(x)g(x).
\end{aligned}$$

Dabei wurde die Stetigkeit von f in x verwendet.

c) *Quotientenregel.* Wir behandeln zunächst den Spezialfall $f = 1$.

$$\begin{aligned}
\left(\frac{1}{g}\right)'(x) &= \lim_{h \rightarrow 0} \frac{1}{h} \left(\frac{1}{g(x+h)} - \frac{1}{g(x)} \right) \\
&= \lim_{h \rightarrow 0} \frac{1}{g(x+h)g(x)} \left(\frac{g(x) - g(x+h)}{h} \right) = \frac{-g'(x)}{g(x)^2}.
\end{aligned}$$

Der allgemeine Fall folgt hieraus mithilfe der Produktregel:

$$\begin{aligned}
\left(\frac{f}{g}\right)'(x) &= \left(f \cdot \frac{1}{g}\right)'(x) = f'(x) \frac{1}{g(x)} + f(x) \frac{-g'(x)}{g(x)^2} \\
&= \frac{f'(x)g(x) - f(x)g'(x)}{g(x)^2}.
\end{aligned}$$

Beispiele

(15.10) Sei $f_n(x) = x^n$, $n \in \mathbb{N}$.

Behauptung: $f'_n(x) = nx^{n-1}$.

Beweis durch vollständige Induktion nach n .

Die Fälle $n = 0, 1, 2$ wurden bereits in den Beispielen (15.1) bis (15.3) behandelt.

Induktionsschritt $n \rightarrow n+1$. Da $f_{n+1} = f_1 f_n$, folgt aus der Produktregel

$$f'_{n+1}(x) = f'_1(x)f_n(x) + f_1(x)f'_n(x) = 1 \cdot x^n + x(nx^{n-1}) = (n+1)x^n.$$

(15.11) $f: \mathbb{R}^* \rightarrow \mathbb{R}$, $f(x) = \frac{1}{x^n}$, $n \in \mathbb{N}$.

Die Quotientenregel liefert sofort

$$f'(x) = \frac{-(nx^{n-1})}{(x^n)^2} = -nx^{-n-1}.$$

Aus (15.10) und (15.11) zusammen folgt, dass

$$\frac{d}{dx}(x^n) = nx^{n-1} \quad \text{für alle } n \in \mathbb{Z}.$$

(Falls $n < 0$, muss $x \neq 0$ vorausgesetzt werden.)

(15.12) Für die Funktion $\tan x = \frac{\sin x}{\cos x}$ erhalten wir aus der Quotientenregel

$$\tan'(x) = \frac{\sin'(x) \cos(x) - \sin(x) \cos'(x)}{\cos^2(x)} = \frac{\cos^2 x + \sin^2 x}{\cos^2 x} = \frac{1}{\cos^2 x}.$$

Satz 3 (Ableitung der Umkehrfunktion). *Sei $I \subset \mathbb{R}$ ein nicht-triviales (d.h. ein aus mehr als einem Punkt bestehendes) Intervall, $f: I \rightarrow \mathbb{R}$ eine stetige, streng monotone Funktion und $g = f^{-1}: J \rightarrow \mathbb{R}$ die Umkehrfunktion, wobei $J = f(I)$.*

Ist f im Punkt $x \in I$ differenzierbar und $f'(x) \neq 0$, so ist g im Punkt $y := f(x)$ differenzierbar und es gilt

$$g'(y) = \frac{1}{f'(x)} = \frac{1}{f'(g(y))}.$$

Beweis. Sei $\eta_v \in J \setminus \{y\}$ irgendeine Folge mit $\lim_{v \rightarrow \infty} \eta_v = y$. Wir setzen $\xi_v := g(\eta_v)$. Da g stetig ist (§12, Satz 1), ist $\lim_{v \rightarrow \infty} \xi_v = x$. Außerdem ist $\xi_v \neq x$ für alle v , da $g: J \rightarrow I$ bijektiv ist. Nun gilt

$$\lim_{v \rightarrow \infty} \frac{g(\eta_v) - g(y)}{\eta_v - y} = \lim_{v \rightarrow \infty} \frac{\xi_v - x}{f(\xi_v) - f(x)} = \lim_{v \rightarrow \infty} \frac{1}{\frac{f(\xi_v) - f(x)}{\xi_v - x}} = \frac{1}{f'(x)}.$$

Also ist $g'(y) = \frac{1}{f'(x)} = \frac{1}{f'(g(y))}$.

Beispiele

(15.13) $\log: \mathbb{R}_+^* \rightarrow \mathbb{R}$ ist die Umkehrfunktion von $\exp: \mathbb{R} \rightarrow \mathbb{R}$. Daher gilt nach dem vorhergehenden Satz

$$\log'(x) = \frac{1}{\exp'(\log x)} = \frac{1}{\exp(\log x)} = \frac{1}{x}.$$

Anwendung. Aus der Ableitung des Logarithmus lässt sich folgende Darstellung für die Zahl e ableiten:

$$e = \lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n.$$

Beweis. Da $\log'(1) = 1$, folgt

$$\lim_{n \rightarrow \infty} n \log \left(1 + \frac{1}{n}\right) = \lim_{n \rightarrow \infty} \frac{\log(1 + \frac{1}{n})}{\frac{1}{n}} = 1.$$

Nun ist $(1 + \frac{1}{n})^n = \exp(n \log(1 + \frac{1}{n}))$, also wegen der Stetigkeit von $\exp$

$$\lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n = \exp(1) = e, \quad \text{q.e.d.}$$

(15.14) $\arcsin: [-1, 1] \rightarrow \mathbb{R}$ ist die Umkehrfunktion von $\sin: [-\frac{\pi}{2}, \frac{\pi}{2}] \rightarrow \mathbb{R}$. Für $x \in]-1, 1[$ gilt:

$$\arcsin'(x) = \frac{1}{\sin'(\arcsin x)} = \frac{1}{\cos(\arcsin x)}.$$

Sei $y := \arcsin x$. Dann ist $\sin y = x$ und $\cos y = +\sqrt{1-x^2}$, da $y \in [-\frac{\pi}{2}, \frac{\pi}{2}]$. Also haben wir

$$\frac{d \arcsin x}{dx} = \frac{1}{\sqrt{1-x^2}} \quad \text{für } -1 < x < 1.$$

(15.15) $\arctan: \mathbb{R} \rightarrow \mathbb{R}$ ist die Umkehrfunktion von $\tan:]-\frac{\pi}{2}, \frac{\pi}{2}[\rightarrow \mathbb{R}$. Also gilt

$$\arctan'(x) = \frac{1}{\tan'(\arctan x)} = \cos^2(\arctan x).$$

Setzen wir $y := \arctan x$, so folgt

$$x^2 = \tan^2 y = \frac{\sin^2 y}{\cos^2 y} = \frac{1 - \cos^2 y}{\cos^2 y} = \frac{1}{\cos^2 y} - 1,$$

also

$$\cos^2 y = \frac{1}{1+x^2}.$$

Deshalb gilt

$$\frac{d \arctan x}{dx} = \frac{1}{1+x^2}.$$

Es ist bemerkenswert, dass die (relativ) komplizierte Funktion $\arctan$ einen so einfachen Differentialquotienten besitzt.

Satz 4 (Kettenregel). Seien $f: V \rightarrow \mathbb{R}$ und $g: W \rightarrow \mathbb{R}$ Funktionen mit $f(V) \subset W$. Die Funktion f sei im Punkt $x \in V$ differenzierbar und g sei in $y := f(x) \in W$ differenzierbar. Dann ist die zusammengesetzte Funktion

$$g \circ f: V \rightarrow \mathbb{R}$$

im Punkt x differenzierbar und es gilt

$$(g \circ f)'(x) = g'(f(x)) f'(x).$$

Beweis. Wir definieren die Funktion $g^*: W \rightarrow \mathbb{R}$ durch

$$g^*(\eta) := \begin{cases} \frac{g(\eta) - g(y)}{\eta - y}, & \text{falls } \eta \neq y, \\ g'(y), & \text{falls } \eta = y. \end{cases}$$

Da g in y differenzierbar ist, gilt

$$\lim_{\eta \rightarrow y} g^*(\eta) = g^*(y) = g'(y).$$

Außerdem gilt für alle $\eta \in W$

$$g(\eta) - g(y) = g^*(\eta)(\eta - y).$$

Damit erhalten wir

$$\begin{aligned} (g \circ f)'(x) &= \lim_{\xi \rightarrow x} \frac{g(f(\xi)) - g(f(x))}{\xi - x} = \lim_{\xi \rightarrow x} \frac{g^*(f(\xi))(f(\xi) - f(x))}{\xi - x} \\ &= \lim_{\xi \rightarrow x} g^*(f(\xi)) \lim_{\xi \rightarrow x} \frac{f(\xi) - f(x)}{\xi - x} \\ &= g'(f(x)) f'(x), \quad \text{q.e.d.} \end{aligned}$$

Beispiele

(15.16) Sei $f: \mathbb{R} \rightarrow \mathbb{R}$ differenzierbar und $F: \mathbb{R} \rightarrow \mathbb{R}$ definiert durch

$$F(x) := f(ax + b), \quad (a, b \in \mathbb{R}).$$

Dann gilt

$$F'(x) = af'(ax + b).$$

(15.17) Sei $a \in \mathbb{R}$ und $f: \mathbb{R}_+^* \rightarrow \mathbb{R}$, $f(x) = x^a$. Da $x^a = \exp(a \log x)$, liefert die Kettenregel

$$\frac{dx^a}{dx} = \exp'(a \log x) \frac{d}{dx}(a \log x) = \exp(a \log x) \frac{a}{x} = x^a \cdot \frac{a}{x} = ax^{a-1}.$$

Somit gilt die in (15.10) und (15.11) für ganze Exponenten bewiesene Formel für beliebige reelle Exponenten.

(15.18) Wir zeigen, dass sich die Quotientenregel auch aus der Kettenregel ableiten lässt. Sei nämlich $g: V \rightarrow \mathbb{R}$ eine in $x \in V$ differenzierbare Funktion, die nirgends den Wert 0 annimmt. Wir setzen $f: \mathbb{R}^* \rightarrow \mathbb{R}$, $f(x) := \frac{1}{x}$. Dann gilt

$$\frac{1}{g} = f \circ g.$$

Nach Beispiel (15.4) ist $f'(x) = -\frac{1}{x^2}$, also

$$\left(\frac{1}{g}\right)'(x) = f'(g(x))g'(x) = -\frac{1}{g(x)^2}g'(x) = \frac{-g'(x)}{g(x)^2}.$$

Aus diesem Spezialfall folgt, wie wir bereits gesehen haben, mithilfe der Produktregel die allgemeine Quotientenregel.

Ableitungen höherer Ordnung

Die Funktion $f: V \rightarrow \mathbb{R}$ sei in V differenzierbar. Falls die Ableitung $f': V \rightarrow \mathbb{R}$ ihrerseits im Punkt $x \in V$ differenzierbar ist, so heißt

$$\frac{d^2 f(x)}{dx^2} := f''(x) := (f')'(x)$$

die *zweite Ableitung* von f in x .

Allgemein definieren wir durch vollständige Induktion: Eine Funktion $f: V \rightarrow \mathbb{R}$ heißt k -mal differenzierbar im Punkt $x \in V$, falls ein $\varepsilon > 0$ existiert, so dass

$$f|V \cap]x - \varepsilon, x + \varepsilon[\rightarrow \mathbb{R}$$

$(k-1)$ -mal differenzierbar in $V \cap]x - \varepsilon, x + \varepsilon[$ ist, und die $(k-1)$ -te Ableitung von f in x differenzierbar ist. Man verwendet folgende Bezeichnungen:

$$f^{(k)}(x) := \frac{d^k f(x)}{dx^k} := \left(\frac{d}{dx}\right)^k f(x) := \frac{d}{dx} \left(\frac{d^{k-1} f(x)}{dx^{k-1}}\right).$$

Die Funktion $f: V \rightarrow \mathbb{R}$ heißt k -mal differenzierbar in V , wenn f in jedem Punkt $x \in V$ k -mal differenzierbar ist. Sie heißt k -mal stetig differenzierbar in V , wenn überdies die k -te Ableitung $f^{(k)}: V \rightarrow \mathbb{R}$ in V stetig ist.

Unter der 0-ten Ableitung einer Funktion versteht man die Funktion selbst.

AUFGABEN

15.1. Man berechne die Ableitungen der folgenden Funktionen:

$$\begin{aligned} f_k: \mathbb{R}_+^* &\rightarrow \mathbb{R}, \quad k = 0, \dots, 5, \\ f_0(x) &:= x^x, \quad f_1(x) := x^{(x^x)}, \quad f_2(x) := (x^x)^x, \\ f_3(x) &:= x^{(x^a)}, \quad f_4(x) := x^{(a^x)}, \quad f_5(x) := a^{(x^x)}. \end{aligned}$$

Dabei sei a eine positive Konstante.

15.2. Man berechne die Ableitung der Funktionen

$$\sinh: \mathbb{R} \rightarrow \mathbb{R},$$

$$\cosh: \mathbb{R} \rightarrow \mathbb{R},$$

$$\tanh := \frac{\sinh}{\cosh}: \mathbb{R} \rightarrow \mathbb{R}.$$

15.3. Man berechne die Ableitungen der folgenden Funktionen (vgl. Aufgabe 12.2):

$$\operatorname{Arsinh}: \mathbb{R} \rightarrow \mathbb{R},$$

$$\operatorname{Arcosh}:]1, \infty[\rightarrow \mathbb{R}.$$

15.4. Man beweise: Die Funktion $\tanh: \mathbb{R} \rightarrow \mathbb{R}$ ist streng monoton wachsend und bildet $\mathbb{R}$ bijektiv auf $] -1, 1[$ ab. Die Umkehrfunktion

$$\operatorname{Artanh}:] -1, 1[\rightarrow \mathbb{R}$$

ist differenzierbar. Man berechne die Ableitung.

15.5. Für $\alpha \in \mathbb{R}$ sei $F_\alpha: \mathbb{R} \rightarrow \mathbb{R}$ definiert durch

$$F_\alpha(x) := \begin{cases} x^\alpha & \text{für } x > 0, \\ 0 & \text{für } x = 0, \\ -|x|^\alpha & \text{für } x < 0. \end{cases}$$

Für welche α ist F_α im Nullpunkt stetig, für welche α differenzierbar?

15.6. Die Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ sei wie folgt definiert:

$$f(x) := \begin{cases} x^2 \sin\left(\frac{1}{x}\right) & \text{für } x \neq 0, \\ 0 & \text{für } x = 0. \end{cases}$$

Man zeige, dass f in jedem Punkt $x \in \mathbb{R}$ differenzierbar ist und berechne die Ableitung. Ist f' stetig?

15.7. Sei $V \subset \mathbb{R}$ und $a \in V$ ein Häufungspunkt von V . Seien

$$f, g: V \rightarrow \mathbb{R}$$

zwei Funktionen mit folgenden Eigenschaften:

- i) f ist in a differenzierbar und $f(a) = 0$.
- ii) g ist in a stetig.

Man beweise: Die Funktion $fg: V \rightarrow \mathbb{R}$ ist in a differenzierbar mit

$$(fg)'(a) = f'(a)g(a).$$

15.8. Man beweise: Die Funktion

$$f: \mathbb{R}_+ \rightarrow \mathbb{R}, \quad x \mapsto f(x) := \begin{cases} x^x & \text{für } x > 0, \\ 1 & \text{für } x = 0, \end{cases}$$

ist im Nullpunkt stetig, aber nicht differenzierbar.

15.9. Sei $V \subset \mathbb{R}$ ein offenes Intervall und seien $f, g: V \rightarrow \mathbb{R}$ zwei in V differenzierbare Funktionen. Die Funktion $F: V \rightarrow \mathbb{R}$ sei definiert durch

$$F := \max(f, g).$$

Man zeige: Die Funktion F ist überall in V differenzierbar mit evtl. Ausnahme der Punkte der Menge

$$A := \{x \in V : f(x) = g(x)\}.$$

Für jedes $x \in A$ existieren die einseitigen Ableitungen $F'_+(x)$ und $F'_-(x)$, und zwar gilt

$$F'_+(x) = \max(f'(x), g'(x)), \quad F'_-(x) = \min(f'(x), g'(x)).$$

15.10. Man beweise: Für alle $x \in \mathbb{R}$ gilt

$$e^x = \lim_{n \rightarrow \infty} \left(1 + \frac{x}{n}\right)^n.$$

15.11. Es seien $f, g: V \rightarrow \mathbb{R}$ in V n -mal differenzierbare Funktionen. Man beweise durch vollständige Induktion nach n die folgenden Beziehungen:

- i) $\frac{d^n}{dx^n}(f(x)g(x)) = \sum_{k=0}^n \binom{n}{k} f^{(n-k)}(x)g^{(k)}(x), \quad (\text{Leibniz}).$
- ii) $f(x)\frac{d^n g(x)}{dx^n} = \sum_{k=0}^n (-1)^k \binom{n}{k} \frac{d^{n-k}}{dx^{n-k}}(f^{(k)}(x)g(x)).$

15.12. Eine Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ heißt *gerade*, wenn $f(-x) = f(x)$ für alle $x \in \mathbb{R}$, und *ungerade*, wenn $f(-x) = -f(x)$ für alle $x \in \mathbb{R}$.

- i) Man zeige: Die Ableitung einer geraden (ungeraden) Funktion ist ungerade (gerade).
- ii) Sei $f: \mathbb{R} \rightarrow \mathbb{R}$ die Polynomfunktion

$$f(x) = a_0 + a_1x + \dots + a_nx^n, \quad (a_k \in \mathbb{R}).$$

Man beweise: f ist genau dann gerade (ungerade), wenn $a_k = 0$ für alle ungeraden (geraden) Indizes k .

§ 16 Lokale Extrema. Mittelwertsatz. Konvexität

Wir kommen jetzt zu den ersten Anwendungen der Differentiation. Viele Eigenschaften einer Funktion spiegeln sich nämlich in ihrer Ableitung wider. So kann das Auftreten von lokalen Extrema, die Monotonie und die Konvexität mithilfe der Ableitung untersucht werden. Aus Schranken für die Ableitung erhält man Abschätzungen für das Wachstum der Funktion.

Definition. Sei $f:]a, b[\rightarrow \mathbb{R}$ eine Funktion. Man sagt, f habe in $x \in]a, b[$ ein *lokales Maximum (Minimum)*, wenn ein $\varepsilon > 0$ existiert, so dass

$$f(x) \geq f(\xi) \text{ (bzw. } f(x) \leq f(\xi)) \quad \text{für alle } \xi \text{ mit } |x - \xi| < \varepsilon.$$

Trifft in der letzten Zeile das Gleichheitszeichen nur für $\xi = x$ zu, so nennt man x ein *strenges* oder *striktes* lokales Maximum (Minimum).

Extremum ist der gemeinsame Oberbegriff für Maximum und Minimum. Anstelle von lokalem Extremum spricht man auch von *relativem* Extremum.

Satz 1. Die Funktion $f:]a, b[\rightarrow \mathbb{R}$ besitze im Punkt $x \in]a, b[$ ein lokales Extremum und sei in x differenzierbar. Dann ist $f'(x) = 0$.

Beweis. f besitze in x ein lokales Maximum. Dann existiert ein $\varepsilon > 0$, so dass $]x - \varepsilon, x + \varepsilon[\subset]a, b[$ und

$$f(\xi) \leq f(x) \quad \text{für alle } \xi \in]x - \varepsilon, x + \varepsilon[.$$

Da f in x differenzierbar ist, gilt

$$f'(x) = \lim_{\xi \rightarrow x} \frac{f(\xi) - f(x)}{\xi - x} = \lim_{\underbrace{\xi \searrow x}_{\leq 0}} \frac{f(\xi) - f(x)}{\xi - x} = \lim_{\underbrace{\xi \nearrow x}_{\geq 0}} \frac{f(\xi) - f(x)}{\xi - x}.$$

Daraus folgt $f'(x) = 0$.

Für ein lokales Minimum ist der Satz analog zu beweisen.

Bemerkungen

a) $f'(x) = 0$ ist nur eine notwendige, aber nicht hinreichende Bedingung für ein lokales Extremum. Für die Funktion $f(x) = x^3$ gilt z.B. $f'(0) = 0$, sie besitzt aber in 0 kein lokales Extremum.

b) Nach §11, Satz 2, nimmt jede in einem *kompakten* Intervall stetige Funktion $f: [a, b] \rightarrow \mathbb{R}$ ihr absolutes Maximum und ihr absolutes Minimum an. Liegt ein Extremum jedoch am Rand, so ist dort nicht notwendig $f'(x) = 0$, wie man z.B. an der Funktion

$$f: [0, 1] \rightarrow \mathbb{R}, \quad x \mapsto x$$

sieht.

Satz 2 (Satz von Rolle). *Sei $a < b$ und $f: [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion mit $f(a) = f(b)$. Die Funktion f sei in $]a, b[$ differenzierbar. Dann existiert ein $\xi \in]a, b[$ mit $f'(\xi) = 0$.*

Der Satz von Rolle sagt insbesondere, dass zwischen zwei Nullstellen einer differenzierbaren Funktion eine Nullstelle der Ableitung liegt.

Beweis. Falls f konstant ist, ist der Satz trivial. Ist f nicht konstant, so gibt es ein $x_0 \in]a, b[$ mit $f(x_0) > f(a)$ oder $f(x_0) < f(a)$. Dann wird das absolute Maximum (bzw. Minimum) der Funktion $f: [a, b] \rightarrow \mathbb{R}$ in einem Punkt $\xi \in]a, b[$ angenommen. Nach Satz 1 ist $f'(\xi) = 0$, q.e.d.

Corollar 1 (Mittelwertsatz der Differentialrechnung).

Sei $a < b$ und $f: [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion, die in $]a, b[$ differenzierbar ist. Dann existiert ein $\xi \in]a, b[$, so dass

$$\frac{f(b) - f(a)}{b - a} = f'(\xi).$$

Geometrisch bedeutet der Mittelwertsatz, dass die Steigung der Sekante durch die Punkte $(a, f(a))$ und $(b, f(b))$ gleich der Steigung der Tangente an den Graphen von f an einer gewissen Zwischenstelle $(\xi, f(\xi))$ ist (Bild 16.1).

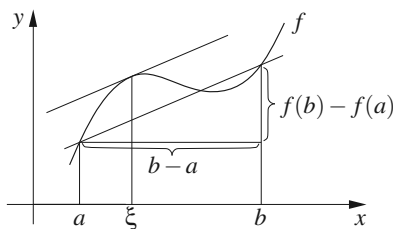


Bild 16.1

Beweis. Wir definieren eine Hilfsfunktion $F: [a, b] \rightarrow \mathbb{R}$ durch

$$F(x) = f(x) - \frac{f(b) - f(a)}{b - a}(x - a).$$

F ist stetig in $[a, b]$ und differenzierbar in $]a, b[$. Da $F(a) = f(a) = F(b)$, existiert nach dem Satz von Rolle ein $\xi \in]a, b[$ mit $F'(\xi) = 0$. Da

$$F'(\xi) = f'(\xi) - \frac{f(b) - f(a)}{b - a},$$

folgt die Behauptung.

Corollar 2. Sei $f: [a, b] \rightarrow \mathbb{R}$ eine stetige, in $]a, b[$ differenzierbare Funktion. Für die Ableitung gelte

$$m \leq f'(\xi) \leq M \quad \text{für alle } \xi \in]a, b[$$

mit gewissen Konstanten $m, M \in \mathbb{R}$. Dann gilt für alle $x_1, x_2 \in [a, b]$ mit $x_1 \leq x_2$ die Abschätzung

$$m(x_2 - x_1) \leq f(x_2) - f(x_1) \leq M(x_2 - x_1).$$

Dies ist eine unmittelbare Folgerung aus dem Mittelwertsatz.

Corollar 3. Sei $f: [a, b] \rightarrow \mathbb{R}$ stetig und in $]a, b[$ differenzierbar mit $f'(x) = 0$ für alle $x \in]a, b[$. Dann ist f konstant.

Dies ist der Fall $m = M = 0$ von Corollar 2.

Als Anwendung geben wir nun eine Charakterisierung der Exponentialfunktion durch ihre Differentialgleichung.

Satz 3. Sei $c \in \mathbb{R}$ eine Konstante und $f: \mathbb{R} \rightarrow \mathbb{R}$ eine differenzierbare Funktion mit

$$f'(x) = cf(x) \quad \text{für alle } x \in \mathbb{R}.$$

Sei $A := f(0)$. Dann gilt

$$f(x) = Ae^{cx} \quad \text{für alle } x \in \mathbb{R}.$$

Beweis. Wir betrachten die Funktion $F(x) := f(x)e^{-cx}$. Nach der Produktregel für die Ableitung ist

$$F'(x) = f'(x)e^{-cx} - cf(x)e^{-cx} = (f'(x) - cf(x))e^{-cx} = 0$$

für alle $x \in \mathbb{R}$, also F konstant. Da $F(0) = f(0) = A$, ist $F(x) = A$ für alle $x \in \mathbb{R}$, woraus folgt

$$f(x) = Ae^{cx} \quad \text{für alle } x \in \mathbb{R}.$$

Bemerkung. Speziell erhält man aus Satz 3: Die Funktion $\exp: \mathbb{R} \rightarrow \mathbb{R}$ ist die eindeutig bestimmte differenzierbare Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ mit $f' = f$ und $f(0) = 1$.

(16.1) Als Beispiel für eine Anwendung von Satz 3 geben wir einen neuen Beweis für die Funktionalgleichung der Exponentialfunktion. Sei $a \in \mathbb{R}$ eine beliebige Konstante und $f: \mathbb{R} \rightarrow \mathbb{R}$ definiert durch

$$f(x) := \exp(a+x).$$

Dann gilt $f'(x) = f(x)$. Aus Satz 3 folgt deshalb $f(x) = f(0)\exp(x)$, d.h.

$$\exp(a+x) = \exp(a)\exp(x).$$

Dies ist aber die Funktionalgleichung der Exponentialfunktion.

Die Funktionen Sinus und Cosinus reproduzieren sich bei zweimaligem Differenzieren bis aufs Vorzeichen. Der folgende Satz sagt, dass dies charakteristisch für diese Funktionen ist.

Satz 4. Sei $f: \mathbb{R} \rightarrow \mathbb{R}$ eine zweimal differenzierbare Funktion mit

$$f''(x) = -f(x) \quad \text{für alle } x \in \mathbb{R}.$$

Dann folgt

$$f(x) = A \cos x + B \sin x \quad \text{mit } A := f(0), B := f'(0).$$

Beweis. Wir betrachten die Funktion

$$g(x) := f(x) - f(0) \cos x - f'(0) \sin x.$$

Für diese Funktion gilt dann ebenso $g'' = -g$, sowie $g(0) = g'(0) = 0$. Sei nun

$$S(x) := g(x)^2 + g'(x)^2.$$

Dann ist

$$S'(x) = 2g(x)g'(x) + 2g'(x)g''(x) = 2g(x)g'(x) - 2g'(x)g(x) = 0,$$

also S konstant. Da $S(0) = 0$, ist S identisch null, also $g(x) = 0$ für alle $x \in \mathbb{R}$, q.e.d.

(16.2) Aus Satz 4 können wir, ähnlich wie in (16.1), die Additions-Theoreme der Funktionen Sinus und Cosinus herleiten. Dazu betrachten wir für konstante $a \in \mathbb{R}$ die Funktionen $g_1, g_2 : \mathbb{R} \rightarrow \mathbb{R}$

$$g_1(x) := \sin(a+x), \quad g_2(x) := \cos(a+x).$$

Es gilt $g_k'' = -g_k$. Aus Satz 4 folgt deshalb $g_k(x) = g_k(0) \cos x + g_k'(0) \sin x$, d.h.

$$\sin(a+x) = \sin a \cos x + \cos a \sin x,$$

$$\cos(a+x) = \cos a \cos x - \sin a \sin x.$$

Dies ist ein bemerkenswert kurzer Beweis der Additions-Theoreme.

Monotonie

Der folgende Satz liefert eine Charakterisierung der Monotonie einer Funktion durch ihre Ableitung.

Satz 5. Sei $f: [a, b] \rightarrow \mathbb{R}$ stetig und in $]a, b[$ differenzierbar.

a) Wenn für alle $x \in]a, b[$ gilt $f'(x) \geq 0$ (bzw. $f'(x) > 0$, $f'(x) \leq 0$, $f'(x) < 0$), so ist f in $[a, b]$ monoton wachsend (bzw. streng monoton wachsend, monoton fallend, streng monoton fallend).

b) Ist f monoton wachsend (bzw. monoton fallend), so folgt $f'(x) \geq 0$ (bzw. $f'(x) \leq 0$) für alle $x \in]a, b[$.

Beweis. a) Wir behandeln nur den Fall, dass $f'(x) > 0$ für alle $x \in]a, b[$ (die übrigen Fälle gehen analog). Angenommen, f sei nicht streng monoton wachsend. Dann gibt es $x_1, x_2 \in [a, b]$ mit $x_1 < x_2$ und $f(x_1) \geq f(x_2)$. Daher existiert nach dem Mittelwertsatz ein $\xi \in]x_1, x_2[$ mit

$$f'(\xi) = \frac{f(x_2) - f(x_1)}{x_2 - x_1} \leq 0.$$

Dies ist ein Widerspruch zur Voraussetzung $f'(\xi) > 0$. Also ist f doch streng monoton wachsend.

b) Sei f monoton wachsend. Dann sind für alle $x, \xi \in]x_1, x_2[, x \neq \xi$, die Differenzenquotienten nicht-negativ:

$$\frac{f(\xi) - f(x)}{\xi - x} \geq 0.$$

Daraus folgt durch Grenzübergang $f'(x) \geq 0$, q.e.d.

Bemerkung. Ist f streng monoton wachsend, so folgt nicht notwendig $f'(x) > 0$ für alle $x \in]x_1, x_2[$, wie das Beispiel der streng monotonen Funktion $f(x) = x^3$ zeigt, deren Ableitung im Nullpunkt verschwindet.

Satz 6. Sei $f:]a, b[\rightarrow \mathbb{R}$ eine differenzierbare Funktion. Im Punkt $x \in]a, b[$ sei f zweimal differenzierbar und es gelte

$$f'(x) = 0 \text{ und } f''(x) > 0 \text{ (bzw. } f''(x) < 0 \text{)}.$$

Dann besitzt f in x ein strenges lokales Minimum (bzw. Maximum).

Bemerkung. Satz 6 gibt nur eine hinreichende, aber nicht notwendige Bedingung für ein strenges Extremum. Die Funktion $f(x) = x^4$ besitzt z.B. für $x = 0$ ein strenges lokales Minimum. Es gilt jedoch $f''(0) = 0$.

Beweis. Sei $f''(x) > 0$. (Der Fall $f''(x) < 0$ ist analog zu beweisen.) Da

$$f''(x) = \lim_{\xi \rightarrow x} \frac{f'(\xi) - f'(x)}{\xi - x} > 0,$$

existiert ein $\varepsilon > 0$, so dass

$$\frac{f'(\xi) - f'(x)}{\xi - x} > 0 \quad \text{für alle } \xi \text{ mit } 0 < |\xi - x| < \varepsilon.$$

Da $f'(x) = 0$, folgt daraus

$$f'(\xi) < 0 \text{ für } x - \varepsilon < \xi < x,$$

$$f'(\xi) > 0 \text{ für } x < \xi < x + \varepsilon.$$

Nach Satz 5 ist deshalb f im Intervall $[x - \varepsilon, x]$ streng monoton fallend und in $[x, x + \varepsilon]$ streng monoton wachsend. f besitzt also in x ein strenges Minimum.

Konvexität

Definition. Sei $D \subset \mathbb{R}$ ein (endliches oder unendliches) Intervall. Eine Funktion $f: D \rightarrow \mathbb{R}$ heißt *konvex*, wenn für alle $x_1, x_2 \in D$ und alle λ mit $0 < \lambda < 1$ gilt

$$f(\lambda x_1 + (1 - \lambda)x_2) \leq \lambda f(x_1) + (1 - \lambda)f(x_2).$$

Die Funktion f heißt *konkav*, wenn $-f$ konvex ist.

Die angegebene Konvexitäts-Bedingung bedeutet (für $x_1 < x_2$), dass der Graph von f im Intervall $[x_1, x_2]$ unterhalb der Sekante durch $(x_1, f(x_1))$ und $(x_2, f(x_2))$ liegt (Bild 16.2).

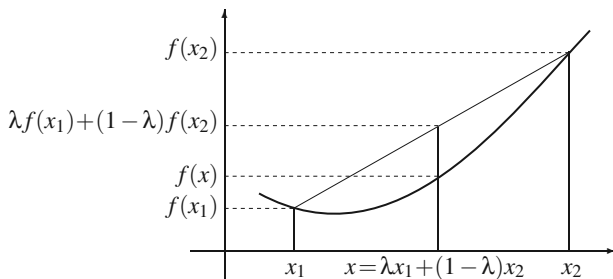


Bild 16.2

Satz 7. Sei $D \subset \mathbb{R}$ ein offenes Intervall und $f: D \rightarrow \mathbb{R}$ eine zweimal differenzierbare Funktion. f ist genau dann konvex, wenn $f''(x) \geq 0$ für alle $x \in D$.

Beweis.

a) Sei zunächst vorausgesetzt, dass $f''(x) \geq 0$ für alle $x \in D$. Dann ist die Ableitung $f': D \rightarrow \mathbb{R}$ nach Satz 5 monoton wachsend. Seien $x_1, x_2 \in D$, $0 < \lambda < 1$ und $x := \lambda x_1 + (1 - \lambda)x_2$. Wir können annehmen, dass $x_1 < x_2$. Dann gilt $x_1 < x < x_2$. Nach dem Mittelwertsatz existieren $\xi_1 \in]x_1, x[$ und $\xi_2 \in]x, x_2[$ mit

$$\frac{f(x) - f(x_1)}{x - x_1} = f'(\xi_1) \leq f'(\xi_2) = \frac{f(x_2) - f(x)}{x_2 - x}.$$

Da $x - x_1 = (1 - \lambda)(x_2 - x_1)$ und $x_2 - x = \lambda(x_2 - x_1)$, folgt daraus

$$\frac{f(x) - f(x_1)}{1 - \lambda} \leq \frac{f(x_2) - f(x)}{\lambda}$$

und weiter

$$f(x) \leq \lambda f(x_1) + (1 - \lambda)f(x_2).$$

Die Funktion f ist also konvex.

b) Sei $f: D \rightarrow \mathbb{R}$ konvex. Angenommen, es gelte nicht $f''(x) \geq 0$ für alle $x \in D$. Dann gibt es ein $x_0 \in D$ mit $f''(x_0) < 0$. Sei $c := f'(x_0)$ und

$$\varphi(x) := f(x) - c(x - x_0) \quad \text{für } x \in D.$$

Dann ist $\varphi: D \rightarrow \mathbb{R}$ eine zweimal differenzierbare Funktion mit $\varphi'(x_0) = 0$ und $\varphi''(x_0) = f''(x_0) < 0$. Nach Satz 6 besitzt φ in x_0 ein strenges lokales Maxi-

mum. Es gibt also ein $h > 0$, so dass $[x_0 - h, x_0 + h] \subset D$ und

$$\varphi(x_0 - h) < \varphi(x_0), \quad \varphi(x_0 + h) < \varphi(x_0).$$

Daraus folgt

$$f(x_0) = \varphi(x_0) > \frac{1}{2}(\varphi(x_0 - h) + \varphi(x_0 + h)) = \frac{1}{2}(f(x_0 - h) + f(x_0 + h)).$$

Setzt man $x_1 := x_0 - h$, $x_2 := x_0 + h$ und $\lambda := \frac{1}{2}$, so ist $x_0 = \lambda x_1 + (1 - \lambda)x_2$, also

$$f(\lambda x_1 + (1 - \lambda)x_2) > \lambda f(x_1) + (1 - \lambda)f(x_2).$$

Dies steht aber im Widerspruch zur Konvexität von f .

Eine einfache Anwendung ist der folgende

Hilfssatz. Seien $p, q \in]1, \infty[$ mit $\frac{1}{p} + \frac{1}{q} = 1$. Dann gilt für alle $x, y \in \mathbb{R}_+$ die Ungleichung

$$x^{1/p} y^{1/q} \leq \frac{x}{p} + \frac{y}{q}.$$

Beweis. Es genügt offenbar, den Hilfssatz für $x, y \in \mathbb{R}_+^*$ zu beweisen. Da für den Logarithmus $\log: \mathbb{R}_+^* \rightarrow \mathbb{R}$ gilt $\log''(x) = -\frac{1}{x^2} < 0$, ist die Funktion $\log$ konkav, also

$$\log\left(\frac{1}{p}x + \frac{1}{q}y\right) \geq \frac{1}{p}\log x + \frac{1}{q}\log y.$$

Nimmt man von beiden Seiten die Exponentialfunktion, so ergibt sich die Behauptung.

p -Norm. Sei p eine reelle Zahl ≥ 1 . Dann definiert man für Vektoren $x = (x_1, \dots, x_n) \in \mathbb{C}^n$ eine Norm $\|x\|_p \in \mathbb{R}_+$ durch

$$\|x\|_p := \left(\sum_{v=1}^n |x_v|^p \right)^{1/p}.$$

Dies ist eine Verallgemeinerung der gewöhnlichen euklidischen Norm, die man für $p = 2$ erhält. Trivial sind folgende beiden Eigenschaften der p -Norm: (i) $\|x\|_p = 0 \Leftrightarrow x = 0$ und (ii) $\|\lambda x\|_p = |\lambda| \cdot \|x\|_p$ für alle $\lambda \in \mathbb{C}$. Zum Beweis der Dreiecksungleichung benötigen wir noch eine andere Ungleichung.

Satz 8 (Höldersche Ungleichung). *Seien $p, q \in]1, \infty[$ mit $\frac{1}{p} + \frac{1}{q} = 1$. Dann gilt für jedes Paar von Vektoren $x = (x_1, \dots, x_n) \in \mathbb{C}^n$, $y = (y_1, \dots, y_n) \in \mathbb{C}^n$*

$$\sum_{v=1}^n |x_v y_v| \leq \|x\|_p \|y\|_q.$$

Beweis. Wir können annehmen, dass $\|x\|_p \neq 0$ und $\|y\|_q \neq 0$, da sonst der Satz trivial ist. Wir setzen

$$\xi_v := \frac{|x_v|^p}{\|x\|_p^p}, \quad \eta_v := \frac{|y_v|^q}{\|y\|_q^q}.$$

Dann ist $\sum_{v=1}^n \xi_v = 1$ und $\sum_{v=1}^n \eta_v = 1$. Der Hilfssatz ergibt angewendet auf ξ_v und η_v

$$\frac{|x_v y_v|}{\|x\|_p \|y\|_q} = \xi_v^{1/p} \eta_v^{1/q} \leq \frac{\xi_v}{p} + \frac{\eta_v}{q}.$$

Durch Summation über v erhält man

$$\frac{1}{\|x\|_p \|y\|_q} \sum_{v=1}^n |x_v y_v| \leq \frac{1}{p} + \frac{1}{q} = 1,$$

also die Behauptung.

Bemerkung. Für $p = q = 2$ erhält man aus der Hölderschen Ungleichung die *Cauchy-Schwarz'sche Ungleichung*

$$|\langle x, y \rangle| \leq \|x\|_2 \|y\|_2 \quad \text{für } x, y \in \mathbb{C}^n.$$

Dabei ist

$$\langle x, y \rangle := \sum_{v=1}^n \bar{x}_v y_v$$

das kanonische Skalarprodukt im $\mathbb{C}^n$.

Satz 9 (Minkowskische Ungleichung). *Sei $p \in [1, \infty[$. Dann gilt für alle $x, y \in \mathbb{C}^n$*

$$\|x + y\|_p \leq \|x\|_p + \|y\|_p.$$

Beweis. Für $p = 1$ folgt der Satz direkt aus der Dreiecksungleichung für komplexe Zahlen. Sei nun $p > 1$ und q definiert durch $\frac{1}{p} + \frac{1}{q} = 1$. Es sei $z \in \mathbb{C}^n$ der Vektor mit den Komponenten

$$z_v := |x_v + y_v|^{p-1}, \quad v = 1, \dots, n.$$

Dann ist $z_v^q = |x_v + y_v|^{q(p-1)} = |x_v + y_v|^p$, also

$$\|z\|_q = \|x + y\|_p^{p/q}.$$

Nach der Hölderschen Ungleichung gilt

$$\sum_v |x_v + y_v| \cdot |z_v| \leq \sum_v |x_v z_v| + \sum_v |y_v z_v| \leq (\|x\|_p + \|y\|_p) \|z\|_q,$$

also nach Definition von z

$$\|x + y\|_p^p \leq (\|x\|_p + \|y\|_p) \|x + y\|_p^{p/q}.$$

Da $p - \frac{p}{q} = 1$, folgt daraus die Behauptung.

Die Regeln von de l'Hospital

Als weitere Anwendung des Mittelwertsatzes leiten wir jetzt einige Formeln her, mit denen man manchmal bequem Grenzwerte berechnen kann.

Lemma. a) Sei $f:]0, a[\rightarrow \mathbb{R}$ eine differenzierbare Funktion mit

$$\lim_{x \searrow 0} f(x) = 0 \quad \text{und} \quad \lim_{x \searrow 0} f'(x) =: c \in \mathbb{R}.$$

Dann gilt $\lim_{x \searrow 0} \frac{f(x)}{x} = c$.

b) Sei $f:]a, \infty[\rightarrow \mathbb{R}$ eine differenzierbare Funktion mit

$$\lim_{x \rightarrow \infty} f'(x) =: c \in \mathbb{R}.$$

Dann gilt $\lim_{x \rightarrow \infty} \frac{f(x)}{x} = c$.

Beweis. Wir beweisen nur Teil b). Der (einfachere) Beweis von Teil a) sei der Leserin überlassen.

Wir behandeln zunächst den Spezialfall $c = 0$.

Wegen $\lim_{x \rightarrow \infty} f'(x) = 0$ gibt es zu vorgegebenem $\varepsilon > 0$ ein $x_0 > \max(a, 0)$ mit $|f'(x)| \leq \varepsilon/2$ für $x \geq x_0$. Aus dem Corollar 2 zu Satz 2 folgt daraus

$$|f(x) - f(x_0)| \leq \frac{\varepsilon}{2} (x - x_0) \quad \text{für alle } x \geq x_0.$$

Für alle $x \geq \max(x_0, 2|f(x_0)|/\varepsilon)$ gilt dann

$$\left| \frac{f(x)}{x} \right| \leq \frac{|f(x) - f(x_0)|}{x} + \frac{|f(x_0)|}{x} \leq \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon,$$

woraus die Behauptung folgt.

Der allgemeine Fall wird durch Betrachtung der Funktion $g(x) := f(x) - cx$ auf den gerade betrachteten Spezialfall zurückgeführt.

Ein *Beispiel* für das Lemma ist die uns schon aus (12.5) bekannte Tatsache, dass

$$\lim_{x \rightarrow \infty} \frac{\log x}{x} = 0.$$

Dies folgt mit dem Lemma daraus, dass $\lim_{x \rightarrow \infty} \log'(x) = \lim_{x \rightarrow \infty} \frac{1}{x} = 0$.

Satz 10 (Regeln von de l'Hospital). *Seien $f, g: I \rightarrow \mathbb{R}$ zwei differenzierbare Funktionen auf dem Intervall $I =]a, b[$, ($-\infty \leq a < b \leq \infty$). Es gelte $g'(x) \neq 0$ für alle $x \in I$ und es existiere der Limes*

$$\lim_{x \nearrow b} \frac{f'(x)}{g'(x)} =: c \in \mathbb{R}.$$

Dann folgt:

- 1) Falls $\lim_{x \nearrow b} g(x) = \lim_{x \nearrow b} f(x) = 0$, ist $g(x) \neq 0$ für alle $x \in I$ und

$$\lim_{x \nearrow b} \frac{f(x)}{g(x)} = c.$$

- 2) Falls $\lim_{x \nearrow b} g(x) = \pm\infty$, ist $g(x) \neq 0$ für $x \geq x_0$, ($a < x_0 < b$), und es gilt ebenfalls

$$\lim_{x \nearrow b} \frac{f(x)}{g(x)} = c.$$

Analoge Aussagen gelten für den Grenzübergang $x \searrow a$.

Beweis. Wir beweisen die Regel 2 durch Zurückführung auf Teil b) des Lemmas. (Regel 1 wird analog mithilfe von Teil a) des Lemmas bewiesen.)

Wir stellen zunächst fest, dass die Abbildung $g: I \rightarrow \mathbb{R}$ injektiv ist, denn gäbe es zwei Punkte $x_1 \neq x_2$ in I mit $g(x_1) = g(x_2)$, so erhielte man mit dem Satz von Rolle eine Nullstelle von g' , was im Widerspruch zur Voraussetzung steht. Es folgt, dass g streng monoton ist und g' das Vorzeichen nicht wechselt. Wir nehmen an, dass g streng monoton wächst (andernfalls gehe man zu $-g$ über). Das Bild von I unter der Abbildung g ist dann das Intervall $J =]A, \infty[$ mit $A = \lim_{x \searrow a} g(x)$. Wir bezeichnen mit $\psi := g^{-1}: J \rightarrow I$ die Umkehrabbildung

und mit F die zusammengesetzte Abbildung

$$F := f \circ \psi: J \rightarrow \mathbb{R}.$$

Für die Ableitung von F gilt nach der Kettenregel und dem Satz über die Ableitung der Umkehrfunktion

$$F'(y) = f'(\psi(y))\psi'(y) = \frac{f'(\psi(y))}{g'(\psi(y))}.$$

und aus der Voraussetzung folgt

$$\lim_{y \rightarrow \infty} F'(y) = \lim_{x \nearrow b} \frac{f'(x)}{g'(x)} = c.$$

Aus dem Lemma folgt deshalb $\lim_{y \rightarrow \infty} \frac{F(y)}{y} = c$. Sei nun $x_n \in I$ eine beliebige Folge mit $\lim x_n = b$. Wir setzen $y_n := g(x_n)$. Dann folgt $\lim y_n = \infty$ und es ist

$$\lim_{n \rightarrow \infty} \frac{f(x_n)}{g(x_n)} = \lim_{n \rightarrow \infty} \frac{f(\psi(y_n))}{y_n} = \lim_{n \rightarrow \infty} \frac{F(y_n)}{y_n} = c, \quad \text{q.e.d.}$$

Beispiele

(16.3) Sei $\alpha > 0$. Nach (12.5) gilt $\lim_{x \rightarrow \infty} (\log x / x^\alpha) = 0$. Dies lässt sich auch mit der 2. Regel von de l'Hospital beweisen: Sei $f(x) := \log x$ und $g(x) = x^\alpha$. Die Voraussetzung $\lim_{x \rightarrow \infty} g(x) = \infty$ ist erfüllt. Nun ist $f'(x) = 1/x$ und $g'(x) = \alpha x^{\alpha-1}$, also

$$\lim_{x \rightarrow \infty} \frac{f'(x)}{g'(x)} = \lim_{x \rightarrow \infty} \frac{1}{\alpha x^\alpha} = 0.$$

Daraus folgt

$$\lim_{x \rightarrow \infty} \frac{\log x}{x^\alpha} = \lim_{x \rightarrow \infty} \frac{f(x)}{g(x)} = 0.$$

(16.4) Manchmal kommt man erst nach Umformungen und mehrmaliger Anwendung der Regeln von de l'Hospital zum Ziel. Sei etwa der Grenzwert

$$\lim_{\substack{x \rightarrow 0 \\ x \neq 0}} \left(\frac{1}{\sin x} - \frac{1}{x} \right)$$

zu untersuchen. Es ist

$$\frac{1}{\sin x} - \frac{1}{x} = \frac{x - \sin x}{x \sin x} = \frac{f(x)}{g(x)}$$

mit $f(x) = x - \sin x$ und $g(x) = x \sin x$. Da

$$\lim_{x \rightarrow 0} f(x) = f(0) = 0 \quad \text{und} \quad \lim_{x \rightarrow 0} g(x) = g(0) = 0,$$

ist also zu untersuchen, ob der Limes

$$\lim_{x \rightarrow 0} \frac{f'(x)}{g'(x)} = \lim_{x \rightarrow 0} \frac{1 - \cos x}{\sin x + x \cos x}$$

existiert. Wegen $\lim_{x \rightarrow 0} f'(x) = f'(0) = 0$ und $\lim_{x \rightarrow 0} g'(x) = g'(0) = 0$ kann man erneut Hospital anwenden. Man berechnet

$$f''(x) = \sin x, \quad g''(x) = 2 \cos x - x \sin x.$$

Da $\lim_{x \rightarrow 0} f''(x) = f''(0) = 0$ und $\lim_{x \rightarrow 0} g''(x) = g''(0) = 2$, ergibt sich insgesamt

$$\lim_{x \rightarrow 0} \frac{f(x)}{g(x)} = \lim_{x \rightarrow 0} \frac{f'(x)}{g'(x)} = \lim_{x \rightarrow 0} \frac{f''(x)}{g''(x)} = \frac{f''(0)}{g''(0)} = \frac{0}{2} = 0.$$

Also haben wir bewiesen

$$\lim_{\substack{x \rightarrow 0 \\ x \neq 0}} \left(\frac{1}{\sin x} - \frac{1}{x} \right) = 0,$$

was bedeutet, dass $\frac{1}{\sin x}$ und $\frac{1}{x}$ für $x \searrow 0$ bzw. $x \nearrow 0$ derart gleichartig gegen $+\infty$ bzw. $-\infty$ gehen, dass ihre Differenz gegen 0 konvergiert.

AUFGABEN

16.1. Man untersuche die Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$,

$$f(x) := x^3 + ax^2 + bx,$$

auf lokale Extrema in Abhängigkeit von den Parametern $a, b \in \mathbb{R}$.

16.2. Man beweise, dass die Funktion

$$f: \mathbb{R}_+ \rightarrow \mathbb{R}, \quad f(x) := x^n e^{-x}, \quad (n > 0),$$

genau ein relatives und absolutes Maximum an der Stelle $x = n$ besitzt.

16.3. Man konstruiere eine zweimal stetig differenzierbare Funktion $f:]-1, 1[\rightarrow \mathbb{R}$ mit folgenden Eigenschaften:

- i) f hat in 0 ein striktes lokales Maximum.
- ii) Es gibt kein $\varepsilon > 0$, so dass f im Intervall $[0, \varepsilon]$ monoton fallend ist.

iii) Es gibt kein $\varepsilon > 0$, so dass f im Intervall $[-\varepsilon, 0]$ monoton wachsend ist.

16.4. Das *Legendresche Polynom* n -ter Ordnung $P_n: \mathbb{R} \rightarrow \mathbb{R}$ ist definiert durch

$$P_n(x) := \frac{1}{2^n n!} \cdot \frac{d^n}{dx^n} [(x^2 - 1)^n].$$

Man beweise:

a) P_n hat genau n paarweise verschiedene Nullstellen im Intervall $] -1, 1[$.

b) P_n genügt der Differentialgleichung

$$(1 - x^2)P_n''(x) - 2xP_n'(x) + n(n+1)P_n(x) = 0$$

(Legendresche Differentialgleichung).

Hinweis. Zum Beweis könnten die Formeln aus 15.11 nützlich sein.

16.5. Man beweise, dass jede in einem offenen Intervall $D \subset \mathbb{R}$ konvexe Funktion $f: D \rightarrow \mathbb{R}$ stetig ist.

16.6. Für $x = (x_1, \dots, x_n) \in \mathbb{C}^n$ sei

$$\|x\|_\infty := \max(|x_1|, \dots, |x_n|).$$

Man beweise $\|x\|_\infty = \lim_{p \rightarrow \infty} \|x\|_p$.

16.7. Sei $f: I \rightarrow \mathbb{R}$ eine im Intervall $I \subset \mathbb{R}$ (nicht notwendig stetig) differenzierbare Funktion. Man zeige: Für die Funktion $f': I \rightarrow \mathbb{R}$ gilt der Zwischenwertsatz, d.h. sind $x_1, x_2 \in I$ und $c \in \mathbb{R}$ mit $f'(x_1) < c < f'(x_2)$, so gibt es eine Stelle $x_0 \in I$ mit $f'(x_0) = c$.

16.8. Sei $I \subset \mathbb{R}$ ein offenes Intervall, $x_0 \in I$ ein Punkt und $f: I \rightarrow \mathbb{R}$ eine stetige Funktion, die in $I \setminus \{x_0\}$ differenzierbar sei. Es existiere der Limes

$$\lim_{\substack{x \rightarrow x_0 \\ x \neq x_0}} f'(x) =: c \in \mathbb{R}.$$

Man beweise, dass f in x_0 differenzierbar ist mit $f'(x_0) = c$.

16.9. Sei $a \in \mathbb{R}$ und $\varepsilon > 0$. Die Funktion

$$f:]a - \varepsilon, a + \varepsilon[\rightarrow \mathbb{R}$$

sei zweimal differenzierbar. Man zeige

$$f''(a) = \lim_{h \rightarrow 0} \frac{f(a+h) - 2f(a) + f(a-h)}{h^2}.$$

16.10. a) Man beweise den *verallgemeinerten Mittelwertsatz*:

Sei $a < b$ und seien $f, g : [a, b] \rightarrow \mathbb{R}$ zwei stetige Funktionen, die in $]a, b[$ differenzierbar sind. Dann existiert ein $\xi \in]a, b[$, so dass

$$(f(b) - f(a))g'(\xi) = (g(b) - g(a))f'(\xi).$$

b) Mithilfe des verallgemeinerten Mittelwertsatzes gebe man einen anderen Beweis der Hospitalschen Regeln (Satz 10).

16.11. Man verallgemeinere die Hospitalschen Regeln (Satz 10) auf den Fall, dass in der Voraussetzung statt $\lim_{x \nearrow b} (f'(x)/g'(x)) = c \in \mathbb{R}$ uneigentliche Konvergenz

$$\lim_{x \nearrow b} \frac{f'(x)}{g'(x)} = \infty$$

vorliegt. Es folgt dann (in beiden Regeln)

$$\lim_{x \nearrow b} \frac{f(x)}{g(x)} = \infty.$$

16.12. Man zeige, dass folgende Limites existieren und berechne sie:

a) $\lim_{x \rightarrow 0} x \cot x,$

b) $\lim_{x \rightarrow 0} \left(\cot x - \frac{1}{x} \right),$

c) $\lim_{x \rightarrow 0} \left(\frac{1}{\sin^2 x} - \frac{1}{x^2} \right).$

d) $\lim_{x \rightarrow \infty} \left(\sqrt{x + \sqrt{x}} - \sqrt{x} \right)$

16.13. Gegeben sei die Funktion $F_a(x) := (2 - a^{1/x})^x$, ($x \in \mathbb{R}_+^*$), wobei $0 < a < 1$ ein Parameter sei. Man untersuche, ob die Grenzwerte

$$\lim_{x \searrow 0} F_a(x) \quad \text{und} \quad \lim_{x \rightarrow \infty} F_a(x)$$

existieren und berechne sie gegebenenfalls.

Hinweis. Man betrachte die Funktion $\log F_a(x)$.

§ 17 Numerische Lösung von Gleichungen

Wir beschäftigen uns jetzt mit der Lösung von Gleichungen $f(x) = 0$, wobei f eine auf einem Intervall vorgegebene Funktion ist. Nicht immer kann man die Lösungen, wie dies etwa bei quadratischen Polynomen der Fall ist, durch einen expliziten Ausdruck angeben. Es sind Näherungsmethoden notwendig, bei denen die Lösungen als Grenzwerte von Folgen dargestellt werden, deren einzelne Glieder berechnet werden können. Für die Brauchbarkeit eines Näherungsverfahrens ist es wichtig, Fehlerabschätzungen zu haben, damit man weiß, wann man bei vorgegebener Fehlerschranke das Verfahren abbrechen darf.

Ein Fixpunktsatz

Es tritt häufig das Problem auf, eine Gleichung der Form $f(x) = x$ lösen zu müssen, wo $f: [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion ist. Hier bietet sich folgendes Näherungsverfahren an. Sei x_0 ein Näherungswert und

$$x_n := f(x_{n-1}) \quad \text{für } n \geq 1.$$

Falls die Folge (x_n) wohldefiniert ist (d.h. jedes x_n wieder im Definitionsbereich von f liegt) und gegen ein $\xi \in [a, b]$ konvergiert, so ist ξ eine Lösung der Gleichung, denn aus der Stetigkeit von f folgt

$$\xi = \lim_{n \rightarrow \infty} x_n = \lim_{n \rightarrow \infty} f(x_{n-1}) = f(\xi).$$

Einen wichtigen Fall, in dem das Verfahren konvergiert, enthält der folgende Satz. Bild 17.1 veranschaulicht das Iterationsverfahren am Graphen von f .

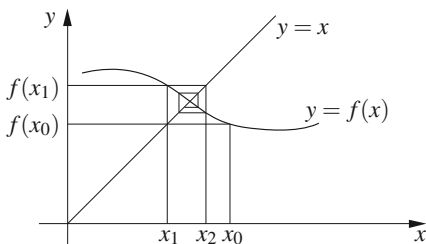


Bild 17.1

Satz 1. Sei $D \subset \mathbb{R}$ ein abgeschlossenes Intervall und $f: D \rightarrow \mathbb{R}$ eine differenzierbare Funktion mit $f(D) \subset D$. Es gebe ein $q < 1$, so dass $|f'(x)| \leq q$ für alle $x \in D$. Sei $x_0 \in D$ beliebig und

$$x_n := f(x_{n-1}) \quad \text{für } n \geq 1.$$

Dann konvergiert die Folge (x_n) gegen die eindeutige Lösung $\xi \in D$ der Gleichung $f(\xi) = \xi$. Es gilt die Fehlerabschätzung

$$|\xi - x_n| \leq \frac{q}{1-q} |x_n - x_{n-1}| \leq \frac{q^n}{1-q} |x_1 - x_0|.$$

Bemerkung. Wie die Fehlerabschätzung zeigt, kann man aus der Differenz zweier aufeinanderfolgender Näherungswerte auf die Genauigkeit der Näherung schließen. Für $q \leq \frac{1}{2}$ etwa ist der Fehler der n -ten Näherung nicht größer als der Unterschied zwischen der $(n-1)$ -ten und der n -ten Näherung.

Das Verfahren konvergiert umso schneller, je kleiner q ist. Dies kann man manchmal durch geeignete Umformungen erreichen. Es sei etwa die Gleichung $F(x) = 0$ zu lösen, wo F eine stetig differenzierbare Funktion ist. Für einen Näherungswert x^* der Lösung sei $F'(x^*) =: c \neq 0$. Setzt man $f(x) := x - \frac{1}{c}F(x)$, so ist die Gleichung $F(\xi) = 0$ äquivalent mit $f(\xi) = \xi$. Es gilt $f'(x^*) = 0$, also ist $|f'(x)|$ klein, falls x hinreichend nahe bei x^* liegt.

Beweis von Satz 1

a) Aus dem Mittelwertsatz erhält man

$$|f(x) - f(y)| \leq q|x - y| \quad \text{für alle } x, y \in D.$$

Daraus folgt insbesondere

$$|x_{n+1} - x_n| = |f(x_n) - f(x_{n-1})| \leq q|x_n - x_{n-1}|$$

und durch Induktion über n

$$|x_{n+1} - x_n| \leq q^n |x_1 - x_0| \quad \text{für alle } n \in \mathbb{N}.$$

Da

$$x_{n+1} = x_0 + \sum_{k=0}^n (x_{k+1} - x_k),$$

und die Reihe $\sum_{k=0}^{\infty} (x_{k+1} - x_k)$ nach dem Majorantenkriterium konvergiert, existiert

$$\xi := \lim_{n \rightarrow \infty} x_n.$$

Weil D ein abgeschlossenes Intervall ist, liegt ξ in D und genügt nach dem eingangs Bemerkten der Gleichung $\xi = f(\xi)$.

b) Zur Eindeutigkeit. Ist η eine weitere Lösung der Gleichung $\eta = f(\eta)$, so gilt

$$|\xi - \eta| = |f(\xi) - f(\eta)| \leq q|\xi - \eta|,$$

woraus wegen $q < 1$ folgt $|\xi - \eta| = 0$, also $\xi = \eta$.

c) Fehlerabschätzung. Für alle $n \geq 1$ und $k \geq 1$ gilt

$$|x_{n+k} - x_{n+k-1}| \leq q^k |x_n - x_{n-1}|.$$

Da $\xi - x_n = \sum_{k=1}^{\infty} (x_{n+k} - x_{n+k-1})$, folgt daraus

$$|\xi - x_n| \leq \sum_{k=1}^{\infty} q^k |x_n - x_{n-1}| = \frac{q}{1-q} |x_n - x_{n-1}| \leq \frac{q^n}{1-q} |x_1 - x_0|.$$

(17.1) Als *Beispiel* wollen wir das Maximum der Funktion $F: \mathbb{R}_+^* \rightarrow \mathbb{R}$,

$$F(x) := \frac{1}{x^5 (e^{1/x} - 1)}$$

bestimmen, vgl. Bild 17.2.

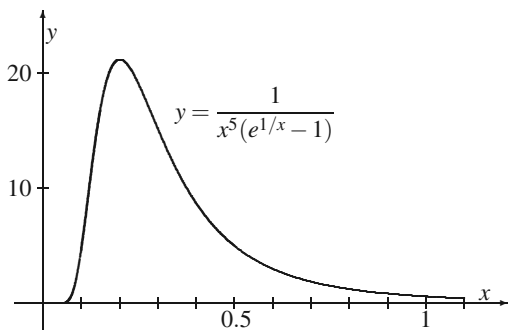


Bild 17.2 Die Plancksche Strahlungsfunktion

Die Funktion F hängt eng mit der *Planckschen Strahlungsfunktion*

$$J(\lambda) = \frac{c^2 h}{\lambda^5 (\exp(\frac{ch}{\lambda k T}) - 1)}$$

zusammen, welche die Strahlungsintensität eines schwarzen Körpers bei der absoluten Temperatur T in Abhängigkeit von der Wellenlänge λ angibt; dabei ist c die Lichtgeschwindigkeit, h die Plancksche und k die Boltzmannsche Konstante. Setzt man $x = \frac{kT}{ch}\lambda$, so ist

$$J(\lambda) = \frac{k^5 T^5}{c^3 h^4} F(x).$$

Für $x > 0$ ist

$$F'(x) = -\frac{5x^4 (e^{1/x} - 1) - x^3 e^{1/x}}{x^{10} (e^{1/x} - 1)^2},$$

also $F'(x) = 0$ genau dann, wenn

$$5x(e^{1/x} - 1) - e^{1/x} = 0.$$

Substituiert man $t := 1/x$, so ist dies äquivalent mit

$$5(1 - e^{-t}) = t.$$

Mit $f(t) := 5(1 - e^{-t})$ hat man also die Gleichung $f(t) = t$ zu lösen. Wir zeigen zunächst, dass die Gleichung in $\mathbb{R}_+$ genau eine Lösung t^* besitzt, die im Intervall $[4, 5]$ liegt. Es ist $f'(t) = 5e^{-t}$, also $f'(t) > 1$ für $t < \log 5$. Im Intervall $[0, \log 5]$ ist also die Funktion $f(t) - t$ streng monoton wachsend. Wegen $f(0) = 0$ gilt $f(t) > t$ für alle $t \in]0, \log 5]$. Für $t > \log 5$ gilt $f'(t) < 1$, also ist die Funktion $f(t) - t$ im Intervall $[\log 5, \infty[$ streng monoton fallend, hat also dort höchstens eine Nullstelle. Wegen

$$f(4) = 4.90\dots > 4,$$

$$f(5) = 4.96\dots < 5$$

gibt es nach dem Zwischenwertsatz tatsächlich eine Nullstelle t^* von $f(t) - t$ im Intervall $[4, 5]$. Es ist

$$q := \sup_{t \in [4, 5]} |f'(t)| = f'(4) = 5e^{-4} = 0.09157\dots,$$

$$\frac{q}{1-q} = 0.1008\dots,$$

also konvergiert die Folge $t_0 := 5$, $t_{n+1} := f(t_n)$, gegen t^* und man hat die Fehlerabschätzung

$$|t^* - t_n| \leq 0.101 |t_n - t_{n-1}|.$$

Man braucht also nur solange zu rechnen, bis die Differenz aufeinander folgender Glieder eine vorgegebene Fehlerschranke ε unterschreitet. Führen wir dies in ARIBAS für $\varepsilon = 10^{-6}$ mit folgender Programmschleife durch

```
==> t := 5.0;
      eps := 10**-6; delta := 1;
      while delta > eps do
        writeln(t);
        t0 := t;
        t := 5*(1 - exp(-t0));
        delta := abs(t-t0);
      end;
      t.
```

so erhalten wir die Ausgabe

```
5.00000000
4.96631026
4.96515593
4.96511569
4.96511428
-: 4.96511423
```

Also ist $t^* = 4.965114\dots$. Für das ursprüngliche Problem bedeutet das, dass die Gleichung $F'(x) = 0$ in $\mathbb{R}_+^*$ genau eine Lösung hat und zwar

$$x^* = \frac{1}{t^*} = 0.2014052 \pm 10^{-7}.$$

Da $\lim_{x \searrow 0} F(x) = 0$ und $\lim_{x \rightarrow \infty} F(x) = 0$, hat die Funktion F an der Stelle x^* ihr einziges Maximum. Die maximale Strahlungsintensität eines schwarzen Körpers der Temperatur T liegt also bei der Wellenlänge

$$\lambda_{\max} = 0.2014 \frac{ch}{kT}.$$

Das Newtonsche Verfahren

Das Newtonsche Verfahren zur Lösung der Gleichung $f(x) = 0$ besteht darin, bei einem Näherungswert x_0 den Graphen von f durch die Tangente zu ersetzen und deren Schnittpunkt mit der x -Achse als neuen Näherungswert x_1 zu benützen und dann das Verfahren zu iterieren, vgl. Bild 17.3.

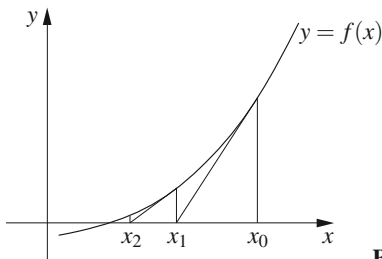


Bild 17.3

Formelmäßig ausgedrückt bedeutet das

$$x_{n+1} := x_n - \frac{f(x_n)}{f'(x_n)}, \quad (n \in \mathbb{N}).$$

Sei f in dem abgeschlossenen Intervall D definiert und stetig differenzierbar mit $f'(x) \neq 0$ für alle $x \in D$. Falls die durch die obige Iterationsvorschrift gebildete Folge (x_n) wohldefiniert ist und gegen ein $\xi \in D$ konvergiert, so folgt aus Stetigkeitsgründen

$$\xi = \xi - \frac{f(\xi)}{f'(\xi)}, \quad \text{also} \quad f(\xi) = 0.$$

Im Allgemeinen braucht das Verfahren jedoch nicht zu konvergieren (Bild 17.4).

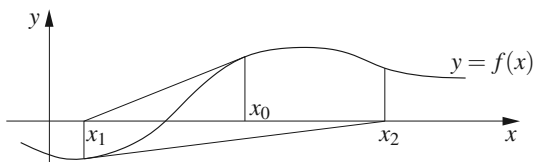


Bild 17.4

Einen wichtigen Fall, in dem Konvergenz auftritt, enthält der folgende Satz.

Satz 2. Es sei $f: [a, b] \rightarrow \mathbb{R}$ eine zweimal differenzierbare konvexe Funktion mit $f(a) < 0$ und $f(b) > 0$. Dann gilt:

- Es gibt genau ein $\xi \in]a, b[$ mit $f(\xi) = 0$.
- Ist $x_0 \in [a, b]$ ein beliebiger Punkt mit $f(x_0) \geq 0$, so ist die Folge

$$x_{n+1} := x_n - \frac{f(x_n)}{f'(x_n)}, \quad (n \in \mathbb{N}),$$

wohldefiniert und konvergiert monoton fallend gegen ξ .

- c) Gilt $f'(\xi) \geq C > 0$ und $f''(x) \leq K$ für alle $x \in]\xi, b[$, so hat man für jedes $n \geq 1$ die Abschätzungen

$$|x_{n+1} - x_n| \leq |\xi - x_n| \leq \frac{K}{2C} |x_n - x_{n-1}|^2.$$

Bemerkungen

1) Analoge Aussagen gelten natürlich auch, falls f konkav ist oder $f(a) > 0$ und $f(b) < 0$ gilt.

2) Die Fehlerabschätzung sagt, dass beim Newtonschen Verfahren sogenannte *quadratische Konvergenz* vorliegt. Ist etwa $\frac{K}{2C}$ größenordnungsmäßig gleich 1 und stimmen x_{n-1} und x_n auf k Dezimalen überein, so ist der Näherungswert x_n auf $2k$ Dezimalstellen genau und bei jedem weiteren Iterationsschritt verdoppelt sich die Zahl der gültigen Stellen.

Beweis von Satz 2

a) Da $f''(x) \geq 0$ für alle $x \in]a, b[$, ist die Funktion f' im ganzen Intervall $[a, b]$ monoton wachsend. Nach §11, Satz 2, existiert ein $q \in [a, b]$ mit

$$f(q) = \inf \{f(x) : x \in [a, b]\} < 0.$$

Falls $q \neq a$, gilt $f'(q) = 0$, also $f'(x) \leq 0$ für $x \leq q$. Die Funktion f ist also im Intervall $[a, q]$ monoton fallend und kann dort keine Nullstelle haben.

In jedem Fall liegen alle Nullstellen von $f: [a, b] \rightarrow \mathbb{R}$ im Intervall $]q, b[$ und nach dem Zwischenwertsatz gibt es dort mindestens eine Nullstelle. Angenommen, es gäbe zwei Nullstellen $\xi_1 < \xi_2$. Nach dem Mittelwertsatz existiert ein $t \in]q, \xi_1[$ mit

$$f'(t) = \frac{f(\xi_1) - f(q)}{\xi_1 - q} = \frac{-f(q)}{\xi_1 - q} > 0,$$

also gilt auch $f'(x) > 0$ für alle $x \geq \xi_1$. Die Funktion f ist also im Intervall $[\xi_1, b]$ streng monoton wachsend und kann keine zweite Nullstelle $\xi_2 > \xi_1$ besitzen.

b) Sei $x_0 \in [a, b]$ mit $f(x_0) \geq 0$. Dann ist notwendig $x_0 \geq \xi$. Wir beweisen durch Induktion, dass für die durch

$$x_{n+1} := x_n - \frac{f(x_n)}{f'(x_n)}$$

definierte Folge gilt $f(x_n) \geq 0$ und $\xi \leq x_n \leq x_{n-1}$ für alle n .

Induktionsschritt $n \rightarrow n+1$. Aus $x_n \geq \xi$ folgt $f'(x_n) \geq f'(\xi) > 0$, also $\frac{f(x_n)}{f'(x_n)} \geq 0$ und daher $x_{n+1} \leq x_n$. Als nächstes zeigen wir $f(x_{n+1}) \geq 0$.

Dazu betrachten wir die Hilfsfunktion

$$\varphi(x) := f(x) - f(x_n) - f'(x_n)(x - x_n).$$

Wegen der Monotonie von f' gilt

$$\varphi'(x) = f'(x) - f'(x_n) \leq 0 \quad \text{für } x \leq x_n.$$

Da $\varphi(x_n) = 0$, ist $\varphi(x) \geq 0$ für $x \leq x_n$, also insbesondere

$$0 \leq \varphi(x_{n+1}) = f(x_{n+1}) - f(x_n) - f'(x_n)(x_{n+1} - x_n) = f(x_{n+1}).$$

Wegen $f(x_{n+1}) \geq 0$ muss aber $x_{n+1} \geq \xi$ gelten, da man sonst einen Widerspruch zum Zwischenwertsatz erhielte.

Wir haben damit bewiesen, dass die Folge (x_n) monoton fällt und durch ξ nach unten beschränkt ist. Also existiert $\lim x_n =: x^*$. Nach dem eingangs Bemerkten gilt dann $f(x^*) = 0$ und wegen der Eindeutigkeit der Nullstelle ist $x^* = \xi$.

c) Da f' monoton wächst und $f'(\xi) \geq C$, gilt $f'(x) \geq C$ für alle $x \geq \xi$. Daraus folgt $f(x) \geq C(x - \xi)$ für alle $x \geq \xi$, insbesondere

$$|\xi - x_n| \leq \frac{f(x_n)}{C}.$$

Um $f(x_n)$ abzuschätzen, betrachten wir die Hilfsfunktion

$$\psi(x) := f(x) - f(x_{n-1}) - f'(x_{n-1})(x - x_{n-1}) - \frac{K}{2}(x - x_{n-1})^2.$$

Differentiation ergibt

$$\begin{aligned} \psi'(x) &= f'(x) - f'(x_{n-1}) - K(x - x_{n-1}), \\ \psi''(x) &= f''(x) - K \leq 0 \quad \text{für alle } x \in]\xi, b[. \end{aligned}$$

Die Funktion ψ' ist also im Intervall $[\xi, b]$ monoton fallend. Da $\psi'(x_{n-1}) = 0$, folgt $\psi'(x) \geq 0$ für $x \in [\xi, x_{n-1}]$. Da auch $\psi(x_{n-1}) = 0$, folgt weiter $\psi(x) \leq 0$ für $x \in [\xi, x_{n-1}]$, insbesondere $\psi(x_n) \leq 0$, d.h.

$$f(x_n) \leq \frac{K}{2}(x_n - x_{n-1})^2, \quad \text{also}$$

$$|\xi - x_n| \leq \frac{f(x_n)}{C} \leq \frac{K}{2C}(x_n - x_{n-1})^2.$$

Damit ist Satz 2 vollständig bewiesen.

(17.2) Beispiel. Sei k eine natürliche Zahl ≥ 2 und $a \in \mathbb{R}_+^*$. Wir betrachten die Funktion

$$f: \mathbb{R}_+ \rightarrow \mathbb{R}, \quad f(x) := x^k - a.$$

Es ist $f'(x) = kx^{k-1}$ und $f''(x) = k(k-1)x^{k-2} \geq 0$ für $x \geq 0$, also f konvex. Das Newtonsche Verfahren zur Nullstellenberechnung ist daher anwendbar. Es gilt

$$x - \frac{f(x)}{f'(x)} = x - \frac{x^k - a}{kx^{k-1}} = \frac{1}{k} \left((k-1)x + \frac{a}{x^{k-1}} \right).$$

Für beliebiges x_0 mit $x_0^k > a$ konvergiert deshalb die Folge

$$x_{n+1} := \frac{1}{k} \left((k-1)x_n + \frac{a}{x_n^{k-1}} \right), \quad (n \in \mathbb{N}),$$

gegen $\sqrt[k]{a}$. (Falls $x_0^k < a$, ist $x_1^k > a$ und das Verfahren konvergiert dann ebenfalls.) Wir haben somit das in §6 beschriebene Verfahren zur Wurzelberechnung als Spezialfall des Newton-Verfahrens wiedergefunden.

AUFGABEN

17.1. Sei $k > 0$ eine natürliche Zahl. Man zeige, dass die Gleichung $x = \tan x$ im Intervall $(k - \frac{1}{2})\pi < x < (k + \frac{1}{2})\pi$ genau eine Lösung ξ besitzt und dass die Folge

$$\begin{aligned} x_0 &:= \left(k + \frac{1}{2}\right)\pi \\ x_{n+1} &:= k\pi + \arctan x_n, \quad (n \in \mathbb{N}), \end{aligned}$$

gegen ξ konvergiert. Man berechne ξ mit einer Genauigkeit von 10^{-6} für die Fälle $k = 1, 2, 3$.

17.2. Man zeige, dass die Gleichung

$$(*) \quad \log(1-x) = -2x$$

im Intervall $0 < x < 1$ genau eine Lösung x_* besitzt und berechne sie mit einer Genauigkeit von 10^{-6} .

Hinweis. Die Gleichung $(*)$ ist äquivalent mit der Fixpunktgleichung

$$x = 1 - e^{-2x}.$$

17.3. Man berechne alle reellen Nullstellen des Polynoms $f(x) = x^5 - x - \frac{1}{5}$ mit einer Genauigkeit von 10^{-6} .

17.4. a) Nach Beispiel (17.2) wird das Newtonsche Verfahren zur Berechnung der 3. Wurzel von $a \in \mathbb{R}_+^*$ durch die Iterationsvorschrift $x_{n+1} = \frac{1}{3}(2x_n + a/x_n^2)$ mit beliebigem Anfangswert $x_0 > 0$ gegeben.

Man zeige, dass auch die durch

$$x_{n+1} = \frac{1}{2}\left(x_n + \frac{a}{x_n^2}\right)$$

rekursiv definierte Folge gegen $\sqrt[3]{a}$ konvergiert und vergleiche die Konvergenzgeschwindigkeit beider Verfahren.

b) Man untersuche das Konvergenzverhalten der durch

$$x_{n+1} = \frac{1}{2}\left(x_n + \frac{a}{x_n^3}\right) \quad \text{bzw.} \quad x_{n+1} = \frac{1}{2}\left(x_n + \frac{a}{x_n^4}\right)$$

rekursiv definierten Folgen.

17.5. Man leite ein weitere hinreichende Bedingung für die Konvergenz des Newton-Verfahrens zur Lösung von $f(x) = 0$ her, indem man auf die Funktion

$$F(x) := x - \frac{f(x)}{f'(x)}$$

den Satz 1 anwende.

17.6. Sei $a > 0$ vorgegeben. Die Folge $(a_n)_{n \in \mathbb{N}}$ werde rekursiv definiert durch

$$a_0 := a \quad \text{und} \quad a_{n+1} := a^{a_n} \quad \text{für alle } n \geq 0.$$

a) Man zeige: Die Folge $(a_n)_{n \in \mathbb{N}}$ konvergiert für $1 \leq a \leq e^{1/e}$ und divergiert für $a > e^{1/e}$.

Hinweis. Ein möglicher Grenzwert ist Fixpunkt der Abbildung $x \mapsto a^x$.

b) Man bestimme den (exakten) Wert von $\lim_{n \rightarrow \infty} a_n$ für $a = e^{1/e}$ und eine numerische Näherung (mit einer Genauigkeit von 10^{-6}) von $\lim_{n \rightarrow \infty} a_n$ für $a = 6/5$.

c) Wie ist das Konvergenzverhalten der Folge für einen Anfangswert $a \in]0, 1[$?

§ 18 Das Riemannsche Integral

Die Integration ist neben der Differentiation die wichtigste Anwendung des Grenzwertbegriffs in der Analysis. Wir definieren das Integral zunächst für Treppenfunktionen, wobei noch keine Grenzwertbetrachtungen nötig sind und der elementargeometrische Flächeninhalt von Rechtecken zugrundeliegt. Das Integral allgemeinerer Funktionen wird dann durch Approximation mittels Treppenfunktionen definiert.

Treppenfunktionen

Für $a, b \in \mathbb{R}$, $a < b$, bezeichne $\mathcal{T}[a, b]$ die Menge aller Treppenfunktionen $\varphi: [a, b] \rightarrow \mathbb{R}$. Wie in §10 definiert, heißt eine Funktion $\varphi: [a, b] \rightarrow \mathbb{R}$ Treppenfunktion, falls es eine Unterteilung

$$a = x_0 < x_1 < \dots < x_n = b$$

des Intervalls $[a, b]$ gibt, so dass φ auf jedem offenen Teilintervall $]x_{k-1}, x_k[$ konstant ist. Die Werte von φ in den Teilpunkten sind beliebig.

Wir zeigen nun, dass $\mathcal{T}[a, b]$ ein Untervektorraum des Vektorraums aller reellen Funktionen $f: [a, b] \rightarrow \mathbb{R}$ ist. Dazu sind folgende Eigenschaften nachzuweisen:

- 1) $0 \in \mathcal{T}[a, b]$,
- 2) $\varphi, \psi \in \mathcal{T}[a, b] \Rightarrow \varphi + \psi \in \mathcal{T}[a, b]$,
- 3) $\varphi \in \mathcal{T}[a, b], \lambda \in \mathbb{R} \Rightarrow \lambda\varphi \in \mathcal{T}[a, b]$.

Die Eigenschaften 1) und 3) sind trivial. Es genügt daher, die Aussage 2) zu beweisen. Die Treppenfunktion φ sei definiert bzgl. der Unterteilung

$$Z: a = x_0 < x_1 < \dots < x_n = b$$

und ψ bzgl. der Unterteilung

$$Z': a = x'_0 < x'_1 < \dots < x'_m = b.$$

Nun sei $a = t_0 < t_1 < \dots < t_k = b$ diejenige Unterteilung von $[a, b]$, die alle Teilpunkte von Z und Z' enthält, d.h.

$$\{t_0, t_1, \dots, t_k\} = \{x_0, x_1, \dots, x_n\} \cup \{x'_0, x'_1, \dots, x'_m\}.$$

Dann sind φ und ψ konstant auf jedem Teilintervall $]t_{j-1}, t_j[$, also ist auch $\varphi + \psi$ auf $]t_{j-1}, t_j[$ konstant. Deshalb gilt $\varphi + \psi \in \mathcal{T}[a, b]$.

Definition (Integral für Treppenfunktionen). Sei $\varphi \in \mathcal{T}[a, b]$ definiert bzgl. der Unterteilung

$$a = x_0 < x_1 < \dots < x_n = b$$

und sei $\varphi|_{[x_{k-1}, x_k]} = c_k$ für $k = 1, \dots, n$. Dann setzt man

$$\int_a^b \varphi(x) dx := \sum_{k=1}^n c_k (x_k - x_{k-1}).$$

Geometrische Deutung. Falls $\varphi(x) \geq 0$ für alle $x \in [a, b]$, kann man $\int_a^b \varphi(x) dx$ als die zwischen der x -Achse und dem Graphen von φ liegende Fläche deuten (schraffierte Fläche in Bild 18.1). Falls φ auf einigen Teilintervallen negativ ist, sind die entsprechenden Flächen negativ in Ansatz zu bringen (Bild 18.2).

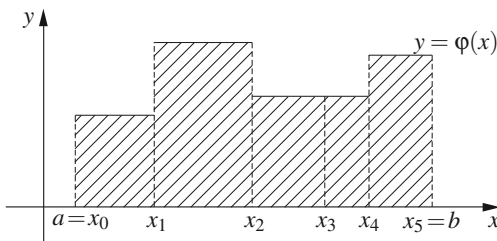


Bild 18.1

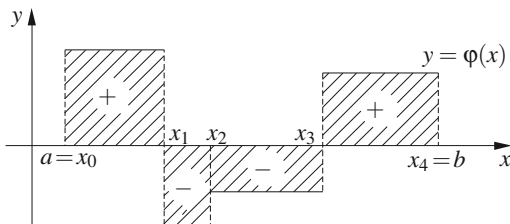


Bild 18.2

Bemerkung. Damit das Integral $\int_a^b \varphi(x) dx$ einer Treppenfunktion wohldefiniert ist, muss man streng genommen noch zeigen, dass die Definition unabhängig von der Unterteilung ist. Es seien

$$Z: a = x_0 < x_1 < \dots < x_n = b,$$

$$Z': a = t_0 < t_1 < \dots < t_m = b$$

zwei Unterteilungen, auf deren offenen Teilintervallen φ konstant ist, und zwar sei

$$\varphi|_{]x_{i-1}, x_i[} = c_i, \quad \varphi|_{]t_{j-1}, t_j[} = c'_j.$$

Wir setzen zur Abkürzung

$$\int_Z \varphi := \sum_{i=1}^n c_i (x_i - x_{i-1}), \quad \int_{Z'} \varphi := \sum_{j=1}^m c'_j (t_j - t_{j-1}).$$

Es ist zu zeigen, dass $\int_Z \varphi = \int_{Z'} \varphi$.

1. Fall. Jeder Teilpunkt von Z sei auch Teilpunkt von Z' , etwa $x_i = t_{k_i}$. Dann gilt

$$x_{i-1} = t_{k_{i-1}} < t_{k_{i-1}+1} < \dots < t_{k_i} = x_i, \quad (1 \leq i \leq n),$$

und $c'_j = c_i$ für $k_{i-1} < j \leq k_i$. Daraus folgt

$$\int_{Z'} \varphi = \sum_{i=1}^n \sum_{j=k_{i-1}+1}^{k_i} c_i (t_j - t_{j-1}) = \sum_{i=1}^n c_i (x_i - x_{i-1}) = \int_Z \varphi.$$

2. Fall. Seien Z und Z' beliebig und sei Z^* die Unterteilung, die alle Teilpunkte von Z und Z' umfasst. Dann gilt nach dem 1. Fall

$$\int_Z \varphi = \int_{Z^*} \varphi = \int_{Z'} \varphi, \quad \text{q.e.d.}$$

Satz 1 (Linearität und Monotonie). *Seien $\varphi, \psi \in \mathcal{T}[a, b]$ und $\lambda \in \mathbb{R}$. Dann gilt:*

$$\text{a) } \int_a^b (\varphi + \psi)(x) dx = \int_a^b \varphi(x) dx + \int_a^b \psi(x) dx.$$

$$\text{b) } \int_a^b (\lambda \varphi)(x) dx = \lambda \int_a^b \varphi(x) dx.$$

$$\text{c) } \varphi \leq \psi \implies \int_a^b \varphi(x) dx \leq \int_a^b \psi(x) dx.$$

Dabei wird für Funktionen $\varphi, \psi: [a, b] \rightarrow \mathbb{R}$ definiert:

$$\varphi \leq \psi \Leftrightarrow \varphi(x) \leq \psi(x) \text{ für alle } x \in [a, b].$$

Bemerkung. Man drückt den Inhalt von Satz 1 auch so aus: Das Integral ist ein lineares, monotones Funktional auf dem Vektorraum $\mathcal{T}[a, b]$.

Beweis. Nach dem oben Bemerkten können φ und ψ bzgl. derselben Unterteilung des Intervalls $[a, b]$ definiert werden. Die Aussagen des Satzes sind dann trivial.

Definition (Oberintegral, Unterintegral). Sei $f: [a, b] \rightarrow \mathbb{R}$ eine beliebige beschränkte Funktion. Dann setzt man

$$\int_a^b f(x) dx := \inf \left\{ \int_a^b \varphi(x) dx : \varphi \in \mathcal{T}[a, b], \varphi \geq f \right\},$$

$$\int_a^b f(x) dx := \sup \left\{ \int_a^b \varphi(x) dx : \varphi \in \mathcal{T}[a, b], \varphi \leq f \right\}.$$

Beispiele

(18.1) Für jede Treppenfunktion $\varphi \in \mathcal{T}[a, b]$ gilt

$$\int_a^b \varphi(x) dx = \int_a^b \varphi(x) dx = \int_a^b \varphi(x) dx.$$

(18.2) Sei $f: [0, 1] \rightarrow \mathbb{R}$ die schon in (10.10) betrachtete Dirichletsche Funktion

$$f(x) := \begin{cases} 1, & \text{falls } x \text{ rational,} \\ 0, & \text{falls } x \text{ irrational.} \end{cases}$$

Dann gilt $\int_0^1 f(x) dx = 1$ und $\int_0^1 f(x) dx = 0$.

Bemerkung. Es gilt stets $\int_a^b f(x) dx \leq \int_a^b f(x) dx$.

Definition. Eine beschränkte Funktion $f: [a, b] \rightarrow \mathbb{R}$ heißt *Riemann-integrierbar*, wenn

$$\int_a^b f(x) dx = \int_a^b f(x) dx.$$

In diesem Fall setzt man

$$\int_a^b f(x) dx := \int_a^b f(x) dx.$$

Bemerkung. Diese Definition des Integrals für Riemann-integrierbare Funktionen $f: [a, b] \rightarrow \mathbb{R}$ ergibt sich zwangsläufig, wenn man das Integral so erklären will, dass es für Treppenfunktionen mit dem schon definierten Integral übereinstimmt und dass aus $f \leq g$ folgt $\int f \leq \int g$. (Hier sei $\int f$ eine Abkürzung für $\int_a^b f(x) dx$, usw.) Denn für jede Treppenfunktion $\varphi \geq f$ gilt dann $\int f \leq \int \varphi$, also $\int f \leq \int^* f$. Ebenso folgt $\int f \geq \int_* f$. Falls also Ober- und Unterintegral von f übereinstimmen, muss der gemeinsame Wert notwendig das Integral von f sein.

(18.3) Beispiele. Nach (18.1) ist jede Treppenfunktion Riemann-integrierbar. Die in (18.2) definierte Funktion ist nicht Riemann-integrierbar.

Schreibweise. Anstelle der Integrationsvariablen x können auch andere Buchstaben verwendet werden (sofern sie nicht mit anderen Bezeichnungen kollidieren):

$$\int_a^b f(x) dx = \int_a^b f(t) dt = \int_a^b f(\xi) d\xi = \dots$$

Satz 2 (Einschließung zwischen Treppenfunktionen). *Eine Funktion $f: [a, b] \rightarrow \mathbb{R}$ ist genau dann Riemann-integrierbar, wenn zu jedem $\varepsilon > 0$ Treppenfunktionen $\varphi, \psi \in \mathcal{T}[a, b]$ existieren mit*

$$\varphi \leq f \leq \psi$$

und

$$\int_a^b \psi(x) dx - \int_a^b \varphi(x) dx \leq \varepsilon.$$

Dies folgt unmittelbar aus der Definition von inf und sup.

Im Folgenden schreiben wir statt Riemann-integrierbar kurz integrierbar.

Satz 3. *Jede stetige Funktion $f: [a, b] \rightarrow \mathbb{R}$ ist integrierbar.*

Beweis. Zu $\varepsilon > 0$ existieren nach §11, Satz 5, Treppenfunktionen $\varphi, \psi \in \mathcal{T}[a, b]$ mit $\varphi \leq f \leq \psi$ und

$$\psi(x) - \varphi(x) \leq \frac{\varepsilon}{b-a} \quad \text{für alle } x \in [a, b].$$

Daher folgt aus Satz 1

$$\int_a^b \psi(x) dx - \int_a^b \varphi(x) dx = \int_a^b (\psi(x) - \varphi(x)) dx \leq \int_a^b \frac{\varepsilon}{b-a} dx = \varepsilon.$$

Nach Satz 2 ist f also integrierbar.

Satz 4. Jede monotone Funktion $f: [a, b] \rightarrow \mathbb{R}$ ist integrierbar.

Beweis. Sei f monoton wachsend (für monoton fallende Funktionen ist der Satz analog zu beweisen). Durch die Punkte

$$x_k := a + k \cdot \frac{b-a}{n}, \quad (k = 0, 1, \dots, n)$$

erhält man eine äquidistante Unterteilung von $[a, b]$. Bezüglich dieser Unterteilung definieren wir Treppenfunktionen $\varphi, \psi \in \mathcal{T}[a, b]$ wie folgt:

$$\varphi(x) := f(x_{k-1}) \quad \text{für } x_{k-1} \leq x < x_k,$$

$$\psi(x) := f(x_k) \quad \text{für } x_{k-1} \leq x < x_k,$$

sowie $\varphi(b) = \psi(b) = f(b)$. Da f monoton wächst, gilt

$$\varphi \leq f \leq \psi$$

und

$$\begin{aligned} \int_a^b \psi(x) dx - \int_a^b \varphi(x) dx &= \sum_{k=1}^n f(x_k) (x_k - x_{k-1}) - \sum_{k=1}^n f(x_{k-1}) (x_k - x_{k-1}) \\ &= \frac{b-a}{n} \left(\sum_{k=1}^n f(x_k) - \sum_{k=1}^n f(x_{k-1}) \right) = \frac{b-a}{n} (f(x_n) - f(x_0)) \leq \varepsilon, \end{aligned}$$

falls n genügend groß ist. Also ist f nach Satz 2 integrierbar.

Satz 5 (Linearität und Monotonie). Seien $f, g: [a, b] \rightarrow \mathbb{R}$ integrierbare Funktionen und $\lambda \in \mathbb{R}$. Dann sind auch die Funktionen $f + g$ und λf integrierbar und es gilt:

$$\text{a) } \int_a^b (f+g)(x) dx = \int_a^b f(x) dx + \int_a^b g(x) dx.$$

$$\text{b) } \int_a^b (\lambda f)(x) dx = \lambda \int_a^b f(x) dx.$$

$$c) f \leq g \implies \int_a^b f(x) dx \leq \int_a^b g(x) dx.$$

Beweis. Wir verwenden das Kriterium von Satz 2.

a) Sei $\varepsilon > 0$ vorgegeben. Dann gibt es nach Voraussetzung Treppenfunktionen $\varphi_1, \psi_1, \varphi_2, \psi_2 \in T[a, b]$ mit

$$\varphi_1 \leq f \leq \psi_1, \quad \varphi_2 \leq g \leq \psi_2$$

und

$$\int_a^b \psi_1(x) dx - \int_a^b \varphi_1(x) dx \leq \frac{\varepsilon}{2} \quad \text{und} \quad \int_a^b \psi_2(x) dx - \int_a^b \varphi_2(x) dx \leq \frac{\varepsilon}{2}.$$

Addition ergibt

$$\varphi_1 + \varphi_2 \leq f + g \leq \psi_1 + \psi_2$$

und

$$\int_a^b (\psi_1(x) + \psi_2(x)) dx - \int_a^b (\varphi_1(x) + \varphi_2(x)) dx \leq \varepsilon.$$

Daraus folgt, dass $f + g$ integrierbar ist und die angegebene Formel gilt.

b) Da die Aussage für $\lambda = 0$ und $\lambda = -1$ trivial ist, genügt es, sie für $\lambda > 0$ zu beweisen. Zu vorgegebenem $\varepsilon > 0$ gibt es Treppenfunktionen φ, ψ mit $\varphi \leq f \leq \psi$ und

$$\int_a^b \psi(x) dx - \int_a^b \varphi(x) dx \leq \frac{\varepsilon}{\lambda}$$

Daraus folgt $\lambda\varphi \leq \lambda f \leq \lambda\psi$ und

$$\int_a^b (\lambda\psi)(x) dx - \int_a^b (\lambda\varphi)(x) dx \leq \varepsilon.$$

Daraus folgt die Behauptung b).

Die Aussage c) ist trivial.

Definition. Für eine Funktion $f: D \rightarrow \mathbb{R}$ definieren wir die Funktionen $f_+, f_-: D \rightarrow \mathbb{R}$ wie folgt:

$$f_+(x) := \begin{cases} f(x), & \text{falls } f(x) > 0, \\ 0 & \text{sonst.} \end{cases}$$

$$f_-(x) := \begin{cases} -f(x), & \text{falls } f(x) < 0, \\ 0 & \text{sonst.} \end{cases}$$

Offenbar gilt $f = f_+ - f_-$ und $|f| = f_+ + f_-$.

Satz 6. Seien $f, g: [a, b] \rightarrow \mathbb{R}$ integrierbare Funktionen. Dann gilt:

a) Die Funktionen f_+ , f_- und $|f|$ sind integrierbar und es gilt

$$\left| \int_a^b f(x) dx \right| \leq \int_a^b |f(x)| dx.$$

b) Für jedes $p \in [1, \infty[$ ist die Funktion $|f|^p$ integrierbar.

c) Die Funktion $fg: [a, b] \rightarrow \mathbb{R}$ ist integrierbar.

Beweis

a) Nach Voraussetzung gibt es zu $\varepsilon > 0$ Treppenfunktionen $\varphi, \psi \in \mathcal{T}[a, b]$ mit $\varphi \leq f \leq \psi$ und

$$\int_a^b (\psi - \varphi)(x) dx \leq \varepsilon.$$

Dann sind auch φ_+ und ψ_+ Treppenfunktionen mit $\varphi_+ \leq f_+ \leq \psi_+$ und

$$\int_a^b (\psi_+ - \varphi_+)(x) dx \leq \int_a^b (\psi - \varphi)(x) dx \leq \varepsilon;$$

also ist f_+ integrierbar. Die Integrierbarkeit von f_- beweist man analog. Nach Satz 5 ist daher auch $|f|$ integrierbar. Die Integral-Abschätzung folgt aus Satz 5c), da $f \leq |f|$ und $-f \leq |f|$.

b) Es genügt, die Integrierbarkeit von $|f|^p$ für den Fall $0 \leq f \leq 1$ zu beweisen. Zu $\varepsilon > 0$ gibt es Treppenfunktionen $\varphi, \psi \in \mathcal{T}[a, b]$ mit

$$0 \leq \varphi \leq f \leq \psi \leq 1$$

und

$$\int_a^b (\psi - \varphi) dx \leq \frac{\varepsilon}{p}.$$

Dann sind auch φ^p und ψ^p Treppenfunktionen mit $\varphi^p \leq f^p \leq \psi^p$ und wegen $\frac{d}{dx}(x^p) = px^{p-1}$ folgt aus dem Mittelwertsatz der Differentialrechnung

$$\psi^p - \varphi^p \leq p(\psi - \varphi).$$

Deshalb ist

$$\int_a^b (\psi^p - \varphi^p)(x) dx \leq p \int_a^b (\psi - \varphi)(x) dx \leq \varepsilon,$$

also f^p integrierbar.

c) Die Behauptung folgt aus Teil b), denn

$$fg = \frac{1}{4} [(f+g)^2 - (f-g)^2].$$

Satz 7 (Mittelwertsatz der Integralrechnung). *Sei $\varphi : [a, b] \rightarrow \mathbb{R}_+$ eine nicht-negative integrierbare Funktion. Dann gibt es zu jeder stetigen Funktion $f : [a, b] \rightarrow \mathbb{R}$ ein $\xi \in [a, b]$, so dass*

$$\int_a^b f(x)\varphi(x) dx = f(\xi) \int_a^b \varphi(x) dx.$$

Im Spezialfall $\varphi = 1$ hat man

$$\int_a^b f(x) dx = f(\xi)(b-a) \quad \text{für ein } \xi \in [a, b].$$

Bemerkung. Geometrisch bedeutet der Mittelwertsatz im Fall $\varphi = 1$ z.B. für eine positive Funktion f , dass die Fläche unter dem Graphen von f gleich der Fläche des Rechtecks mit den Seitenlängen $b-a$ und $f(\xi)$ ist, vgl. Bild 18.3.

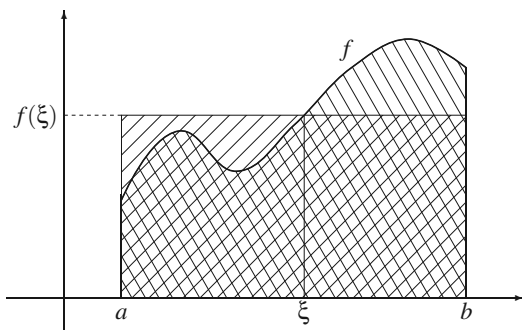


Bild 18.3

Für eine beliebige integrierbare Funktion $f : [a, b] \rightarrow \mathbb{R}$ nennt man

$$M(f) := \frac{1}{b-a} \int_a^b f(x) dx$$

den *Mittelwert* von f über dem Intervall $[a, b]$. Allgemeiner heißt

$$M_\varphi(f) := \frac{1}{\int_a^b \varphi(x) dx} \int_a^b f(x)\varphi(x) dx$$

der (bzgl. φ) *gewichtete Mittelwert* von f (falls $\int_a^b \varphi(x) dx \neq 0$).

Für eine stetige Funktion $f: [a, b] \rightarrow \mathbb{R}$ ist also der Mittelwert gleich dem Wert von f an einer gewissen Zwischenstelle $\xi \in [a, b]$.

Beweis von Satz 7. Nach Satz 6 ist die Funktion $f\varphi$ wieder integrierbar. Wir setzen

$$m := \inf \{f(x) : x \in [a, b]\}, \\ M := \sup \{f(x) : x \in [a, b]\}.$$

Dann gilt $m\varphi \leq f\varphi \leq M\varphi$, also nach Satz 5

$$m \int_a^b \varphi(x) dx \leq \int_a^b f(x)\varphi(x) dx \leq M \int_a^b \varphi(x) dx.$$

Daher existiert ein $\mu \in [m, M]$ mit

$$\int_a^b f(x)\varphi(x) dx = \mu \int_a^b \varphi(x) dx.$$

Nach dem Zwischenwertsatz existiert ein $\xi \in [a, b]$ mit $f(\xi) = \mu$. Daraus folgt die Behauptung.

Riemannsche Summen

Sei $f: [a, b] \rightarrow \mathbb{R}$ eine Funktion,

$$a = x_0 < x_1 < \dots < x_n = b$$

eine Unterteilung von $[a, b]$ und ξ_k ein beliebiger Punkt („Stützstelle“) aus dem Intervall $[x_{k-1}, x_k]$. Das Symbol

$$\mathcal{Z} := ((x_k)_{0 \leq k \leq n}, (\xi_k)_{1 \leq k \leq n})$$

bezeichne die Zusammenfassung der Teilpunkte und der Stützstellen.

Dann heißt

$$S(\mathcal{Z}, f) := \sum_{k=1}^n f(\xi_k) (x_k - x_{k-1})$$

Riemannsche Summe der Funktion f bzgl. $\mathcal{Z}$. Die Riemannsche Summe ist nichts anderes als das Integral einer Treppenfunktion, die die Funktion f an den Stellen ξ_k „interpoliert“, siehe Bild 18.4.

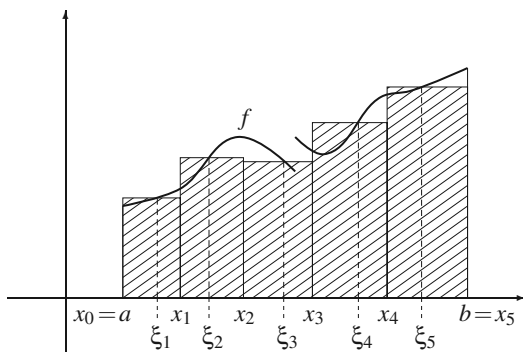


Bild 18.4

Die *Feinheit* (oder *Maschenweite*) von Z ist definiert als

$$\mu(Z) := \max_{1 \leq k \leq n} (x_k - x_{k-1}).$$

Der nächste Satz sagt, dass die Riemannschen Summen einer integrierbaren Funktion gegen das Integral konvergieren, wenn die Feinheit der Unterteilungen gegen Null konvergiert.

Satz 8. Sei $f: [a, b] \rightarrow \mathbb{R}$ eine Riemann-integrierbare Funktion. Dann existiert zu jedem $\varepsilon > 0$ ein $\delta > 0$, so dass für jede Wahl Z von Teilpunkten und Stützstellen der Feinheit $\mu(Z) \leq \delta$ gilt

$$\left| \int_a^b f(x) dx - S(Z, f) \right| \leq \varepsilon.$$

Man kann dies auch so schreiben:

$$\lim_{\mu(Z) \rightarrow 0} S(Z, f) = \int_a^b f(x) dx.$$

Beweis. Sind φ, ψ Treppenfunktionen mit $\varphi \leq f \leq \psi$, so gilt offenbar für alle Zerlegungen Z

$$S(Z, \varphi) \leq S(Z, f) \leq S(Z, \psi).$$

Daraus folgt, dass es genügt, den Satz für den Fall zu beweisen, dass f eine Treppenfunktion ist. Sei f bzgl. der Unterteilung

$$a = t_0 < t_1 < \dots < t_m = b$$

definiert. Da f beschränkt ist, existiert

$$M := \sup \{|f(x)| : x \in [a, b]\} \in \mathbb{R}_+.$$

Sei $Z := ((x_k)_{0 \leq k \leq n}, (\xi_k)_{1 \leq k \leq n})$ irgend eine Unterteilung mit Stützstellen des Intervalls $[a, b]$ und $F \in \mathcal{T}[a, b]$ die durch $F(a) = f(a)$ und

$$F(x) = f(\xi_k) \quad \text{für } x_{k-1} < x \leq x_k \quad (1 \leq k \leq n)$$

definierte Treppenfunktion. Dann gilt

$$S(Z, f) = \int_a^b F(x) dx,$$

also

$$\left| \int_a^b f(x) dx - S(Z, f) \right| \leq \int_a^b |f(x) - F(x)| dx.$$

Die Funktionen f und F stimmen auf allen Teilintervallen $]x_{k-1}, x_k[$ überein, für die $[x_{k-1}, x_k]$ keinen Teilpunkt t_j enthält. Daraus folgt, dass $|f(x) - F(x)|$ auf höchstens $2m$ Teilintervallen $]x_{k-1}, x_k[$ der Gesamtlänge $2m\mu(Z)$ von 0 verschieden sein kann. In jedem Fall gilt aber $|f(x) - F(x)| \leq 2M$, also ist

$$\int_a^b |f(x) - F(x)| dx \leq 4mM\mu(Z).$$

Da dies für $\mu(Z) \rightarrow 0$ gegen 0 konvergiert, folgt die Behauptung des Satzes.

Beispiele

(18.4) Wir berechnen das Integral $\int_0^a x dx$, ($a > 0$),

mittels Riemannscher Summen. Für eine ganze Zahl $n \geq 1$ erhält man durch

$$x_k := \frac{ka}{n}, \quad k = 0, 1, \dots, n,$$

eine äquidistante Unterteilung von $[0, a]$ der Feinheit $\frac{a}{n}$. Als Stützstellen wählen wir $\xi_k = x_k$. Die zugehörige Riemannsche Summe ist dann

$$\begin{aligned} S_n &= \sum_{k=1}^n \frac{ka}{n} \cdot \frac{a}{n} = \frac{a^2}{n^2} \sum_{k=1}^n k \\ &= \frac{a^2}{n^2} \cdot \frac{n(n+1)}{2} = \frac{a^2}{2} \left(1 + \frac{1}{n} \right). \end{aligned}$$

Also folgt

$$\int_0^a x dx = \lim_{n \rightarrow \infty} S_n = \frac{a^2}{2},$$

was die Fläche eines Dreiecks mit Grundlinie a und Höhe a darstellt, vgl. Bild 18.5.

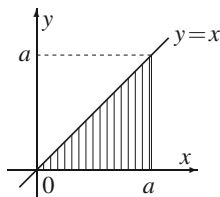


Bild 18.5

Für das nächste Beispiel zu Satz 8 benötigen wir den folgenden Hilfssatz.

Hilfssatz. Sei $t \in \mathbb{R}$ kein ganzzahliges Vielfaches von 2π . Dann gilt für jede natürliche Zahl n

$$\frac{1}{2} + \sum_{k=1}^n \cos kt = \frac{\sin\left(n + \frac{1}{2}\right)t}{2 \sin \frac{1}{2}t}.$$

Beweis. Es gilt $\cos kt = \frac{1}{2}(e^{ikt} + e^{-ikt})$, also

$$\frac{1}{2} + \sum_{k=1}^n \cos kt = \frac{1}{2} \sum_{k=-n}^n e^{ikt}.$$

Nun ist nach der Summenformel für die geometrische Reihe

$$\begin{aligned} \sum_{k=-n}^n e^{ikt} &= e^{-int} \sum_{k=0}^{2n} e^{ikt} = e^{-int} \frac{1 - e^{(2n+1)it}}{1 - e^{it}} \\ &= \frac{e^{i(n+1/2)t} - e^{-i(n+1/2)t}}{e^{it/2} - e^{-it/2}} = \frac{\sin\left(n + \frac{1}{2}\right)t}{\sin \frac{1}{2}t}. \end{aligned}$$

Daraus folgt die Behauptung.

(18.5) Berechnung des Integrals $\int_0^a \cos x dx$, ($a > 0$),

mittels Riemannscher Summen. Wie im Beispiel (18.4) erhalten wir für eine natürliche Zahl $n \geq 1$ durch

$$x_k := \frac{ka}{n}, \quad k = 0, 1, \dots, n,$$

eine äquidistante Unterteilung von $[0, a]$ der Feinheit $\frac{a}{n}$. Als Stützstellen wählen wir $\xi_k = x_k$. Die zugehörige Riemannsche Summe ist dann

$$S_n = \sum_{k=1}^n \frac{a}{n} \cos \frac{ka}{n} = \frac{a}{n} \left(\frac{\sin\left(n + \frac{1}{2}\right)\frac{a}{n}}{2 \sin \frac{a}{2n}} - \frac{1}{2} \right)$$

$$= \frac{\frac{a}{2n}}{\sin \frac{a}{2n}} \cdot \sin \left(a + \frac{a}{2n} \right) - \frac{a}{2n}.$$

Da $\lim_{n \rightarrow \infty} \frac{\sin \frac{a}{2n}}{\frac{a}{2n}} = 1$ (nach §14, Corollar zu Satz 5), folgt

$$\int_0^a \cos x \, dx = \lim_{n \rightarrow \infty} S_n = \sin a.$$

(18.6) Mithilfe von Satz 8 lassen sich die Minkowskische und Höldersche Ungleichung aus §16 auf Integrale verallgemeinern. Sei $f: [a, b] \rightarrow \mathbb{R}$ eine integrierbare Funktion und $p \geq 1$ eine reelle Zahl. Dann definiert man

$$\|f\|_p := \left(\int_a^b |f(x)|^p \, dx \right)^{1/p}.$$

Für integrierbare Funktionen $f, g: [a, b] \rightarrow \mathbb{R}$ gilt dann

a) $\|f + g\|_p \leq \|f\|_p + \|g\|_p$ für alle $p \geq 1$.

b) $\int_a^b |f(x)g(x)| \, dx \leq \|f\|_p \|g\|_q$ für $p, q > 1$ mit $\frac{1}{p} + \frac{1}{q} = 1$.

Satz 9. Sei $a < b < c$ und $f: [a, c] \rightarrow \mathbb{R}$ eine Funktion. f ist genau dann integrierbar, wenn sowohl $f|_{[a, b]}$ also auch $f|_{[b, c]}$ integrierbar sind und es gilt dann

$$\int_a^c f(x) \, dx = \int_a^b f(x) \, dx + \int_b^c f(x) \, dx.$$

Der einfache Beweis sei dem Leser überlassen.

Definition. Man setzt

$$\int_a^a f(x) \, dx := 0,$$

$$\int_a^b f(x) \, dx := - \int_b^a f(x) \, dx, \quad \text{falls } b < a.$$

Bemerkung. Die Formel von Satz 9 gilt nun für beliebige gegenseitige Lage von a, b, c , falls f in $[\min(a, b, c), \max(a, b, c)]$ integrierbar ist.

AUFGABEN

18.1. Man berechne das Integral $\int_0^a x^k dx$, ($k \in \mathbb{N}$, $a \in \mathbb{R}_+^*$)

mittels Riemannscher Summen. Dabei benutze man eine äquidistante Teilung des Intervalls $[0, a]$ und das Ergebnis von Aufgabe 1.5.

18.2. Man berechne das Integral $\int_0^a e^x dx$ mittels Riemannscher Summen ($a > 0$).

18.3. Sei $a > 1$. Man beweise mittels Riemannscher Summen

$$\int_1^a \frac{dx}{x} = \log a.$$

Anleitung. Man wähle folgende Unterteilung:

$$1 = x_0 < x_1 < \dots < x_n = a, \quad \text{wobei} \quad x_k := a^{k/n}.$$

Als Stützstellen wähle man $\xi_k := x_{k-1}$.

18.4. Man beweise

$$\lim_{N \rightarrow \infty} \sum_{n=1}^N \frac{1}{N+n} = \int_1^2 \frac{dx}{x} = \log 2.$$

Bemerkung. Zusammen mit Aufgabe 1.7 folgt daraus die Summe der alternierenden harmonischen Reihe

$$\sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{n} = \log 2.$$

18.5. Seien $f, g: [a, b] \rightarrow \mathbb{R}$ beschränkte Funktionen. Man zeige:

- a) $\int_a^b (f+g)(x) dx \leq \int_a^b f(x) dx + \int_a^b g(x) dx$, (Subadditivität).
 b) $\int_a^b (\lambda f)(x) dx = \lambda \int_a^b f(x) dx$ für alle $\lambda \in \mathbb{R}_+$.

Man gebe ein Beispiel an, für das in a) das Gleichheitszeichen nicht gilt.

18.6. Sei $f: [a, b] \rightarrow \mathbb{R}$ eine beschränkte Funktion, die nur endlich viele Unstetigkeitsstellen hat. Man zeige, dass f Riemann-integrierbar ist.

§ 19 Integration und Differentiation

Während wir im vorigen Paragraphen das Integral in Anlehnung an seine anschauliche Bedeutung als Flächeninhalt definiert haben, zeigen wir hier, dass die Integration die Umkehrung der Differentiation ist, was in vielen Fällen die Möglichkeit zur Berechnung des Integrals liefert.

Für den ganzen Paragraphen sei $I \subset \mathbb{R}$ ein aus mindestens zwei Punkten bestehendes offenes, halboffenes oder abgeschlossenes endliches oder unendliches Intervall.

Unbestimmtes Integral, Stammfunktionen

Während wir bisher Funktionen immer über ein festes abgeschlossenes Intervall integriert haben, betrachten wir jetzt die eine Integrationsgrenze als variabel und erhalten so eine neue Funktion, das „unbestimmte Integral“.

Satz 1. Sei $f: I \rightarrow \mathbb{R}$ eine stetige Funktion und $a \in I$. Für $x \in I$ sei

$$F(x) := \int_a^x f(t) dt.$$

Dann ist die Funktion $F: I \rightarrow \mathbb{R}$ differenzierbar und es gilt $F' = f$.

Beweis. Für $h \neq 0$ ist

$$\frac{F(x+h) - F(x)}{h} = \frac{1}{h} \left(\int_a^{x+h} f(t) dt - \int_a^x f(t) dt \right) = \frac{1}{h} \int_x^{x+h} f(t) dt.$$

Nach dem Mittelwertsatz der Integralrechnung (§18, Satz 7) existiert ein $\xi_h \in [x, x+h]$ (bzw. $\xi_h \in [x+h, x]$, falls $h < 0$) mit

$$\int_x^{x+h} f(t) dt = hf(\xi_h).$$

Da $\lim_{h \rightarrow 0} \xi_h = x$ und f stetig ist, folgt

$$F'(x) = \lim_{h \rightarrow 0} \frac{1}{h} \int_x^{x+h} f(t) dt = \lim_{h \rightarrow 0} \frac{1}{h} (hf(\xi_h)) = f(x).$$

Definition. Eine differenzierbare Funktion $F: I \rightarrow \mathbb{R}$ heißt *Stammfunktion* (oder *primitive Funktion*) einer Funktion $f: I \rightarrow \mathbb{R}$, falls $F' = f$.

Bemerkung. Satz 1 bedeutet, dass das unbestimmte Integral eine Stammfunktion des Integranden ist.

Satz 2. Sei $F: I \rightarrow \mathbb{R}$ eine Stammfunktion von $f: I \rightarrow \mathbb{R}$. Eine weitere Funktion $G: I \rightarrow \mathbb{R}$ ist genau dann Stammfunktion von f , wenn $F - G$ eine Konstante ist.

Beweis.

a) Sei $F - G = c$ mit der Konstanten $c \in \mathbb{R}$. Dann ist $G' = (F - c)' = F' = f$.

b) Sei G Stammfunktion von f , also $G' = f = F'$. Dann gilt $(F - G)' = 0$, daher ist $F - G$ konstant (§16, Corollar 3 zu Satz 2).

Satz 3 (Fundamentalsatz der Differential- und Integralrechnung).

Sei $f: I \rightarrow \mathbb{R}$ eine stetige Funktion und F eine Stammfunktion von f . Dann gilt für alle $a, b \in I$

$$\int_a^b f(x) dx = F(b) - F(a).$$

Beweis. Für $x \in I$ sei

$$F_0(x) := \int_a^x f(t) dt.$$

Ist nun F eine beliebige Stammfunktion von f , so gibt es nach Satz 2 ein $c \in \mathbb{R}$ mit $F - F_0 = c$. Deshalb ist

$$F(b) - F(a) = F_0(b) - F_0(a) = F_0(b) = \int_a^b f(t) dt, \quad \text{q.e.d.}$$

Bezeichnung. Man setzt

$$F(x) \Big|_a^b := F(b) - F(a).$$

Die Formel von Satz 3 schreibt sich dann als

$$\int_a^b f(x) dx = F(x) \Big|_a^b.$$

Hierfür schreibt man abkürzend

$$\int f(x) dx = F(x).$$

Diese Schreibweise ist jedoch insofern problematisch, als F nur bis auf eine Konstante eindeutig bestimmt ist.

Beispiele

Aufgrund von Satz 3 erhält man aus jeder Differentiationsformel eine Formel über Integration. Wir stellen einige Beispiele zusammen.

(19.1) Sei $s \in \mathbb{R}$, $s \neq -1$. Dann gilt

$$\int_a^b x^s dx = \frac{x^{s+1}}{s+1} \Big|_a^b.$$

Dabei ist das Integrationsintervall folgenden Einschränkungen unterworfen: Für $s \in \mathbb{N}$ sind $a, b \in \mathbb{R}$ beliebig; ist s eine ganze Zahl ≤ -2 , so darf 0 nicht im Integrationsintervall liegen; ist s nicht ganz, so ist $[a, b] \subset \mathbb{R}_+^*$ vorauszusetzen (bzw. $[b, a] \subset \mathbb{R}_+^*$, falls $b < a$).

(19.2) Für $a, b > 0$ gilt

$$\int_a^b \frac{dx}{x} = \log x \Big|_a^b.$$

Für $a, b < 0$ gilt

$$\int_a^b \frac{dx}{x} = \log(-x) \Big|_a^b, \quad \text{da} \quad \frac{d}{dx} \log(-x) = \frac{1}{x} \quad \text{für } x < 0.$$

Man kann die beiden Fälle so zusammenfassen:

$$\int \frac{dx}{x} = \log |x| \quad \text{für } x \neq 0.$$

Dabei soll $x \neq 0$ bedeuten: Der Punkt 0 liegt nicht im Integrationsintervall.

$$\textbf{(19.3)} \quad \int \sin x dx = -\cos x.$$

$$\textbf{(19.4)} \quad \int \cos x dx = \sin x.$$

Damit haben wir auf mühelose Weise das in (18.5) mittels Riemannscher Summen berechnete Integral wiedererhalten.

$$(19.5) \quad \int \exp x \, dx = \exp x.$$

$$(19.6) \quad \int \frac{dx}{\sqrt{1-x^2}} = \arcsin x \quad \text{für } |x| < 1.$$

$$(19.7) \quad \int \frac{dx}{1+x^2} = \arctan x.$$

$$(19.8) \quad \int \frac{dx}{\cos^2 x} = \tan x;$$

dabei muss im Integrationsintervall $\cos x \neq 0$ sein.

Die Substitutionsregel

Ein wichtiges Hilfsmittel zur Auswertung von Integralen besteht darin, eine Transformation (Substitution) der Integrationsvariablen durchzuführen. Durch geschickte Wahl der Substitution kann man oft das Integral vereinfachen und zugänglicher machen.

Satz 4 (Substitutionsregel). *Sei $f: I \rightarrow \mathbb{R}$ eine stetige Funktion und $\varphi: [a, b] \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion mit $\varphi([a, b]) \subset I$. Dann gilt*

$$\int_a^b f(\varphi(t)) \varphi'(t) \, dt = \int_{\varphi(a)}^{\varphi(b)} f(x) \, dx.$$

Beweis. Sei $F: I \rightarrow \mathbb{R}$ eine Stammfunktion von f . Für die Funktion $F \circ \varphi: [a, b] \rightarrow \mathbb{R}$ gilt nach der Kettenregel

$$(F \circ \varphi)'(t) = F'(\varphi(t)) \varphi'(t) = f(\varphi(t)) \varphi'(t).$$

Daraus folgt nach Satz 3

$$\int_a^b f(\varphi(t)) \varphi'(t) \, dt = (F \circ \varphi)(t) \Big|_a^b = F(\varphi(b)) - F(\varphi(a)) = \int_{\varphi(a)}^{\varphi(b)} f(x) \, dx.$$

Bezeichnung. Unter Verwendung der symbolischen Schreibweise

$$d\varphi(t) := \varphi'(t)dt$$

lautet die Substitutionsregel

$$\int_a^b f(\varphi(t)) d\varphi(t) = \int_{\varphi(a)}^{\varphi(b)} f(x) dx.$$

In dieser Form ist sie besonders einfach zu merken, denn man hat einfach x durch $\varphi(t)$ zu ersetzen. Läuft t von a nach b , so läuft $x = \varphi(t)$ von $\varphi(a)$ nach $\varphi(b)$.

Beispiele

$$(19.9) \quad \int_a^b f(t+c) dt = \int_{a+c}^{b+c} f(x) dx, \quad (\text{Substitution } \varphi(t) = t+c).$$

$$(19.10) \quad \text{Für } c \neq 0 \text{ gilt } \int_a^b f(ct) dt = \frac{1}{c} \int_{ac}^{bc} f(x) dx, \quad (\varphi(t) = ct).$$

$$(19.11) \quad \int_a^b t f(t^2) dt = \frac{1}{2} \int_{a^2}^{b^2} f(x) dx, \quad (\varphi(t) = t^2).$$

(19.12) Sei $\varphi: [a, b] \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion mit $\varphi(t) \neq 0$ für alle $t \in [a, b]$. Dann gilt nach (19.2)

$$\int_a^b \frac{\varphi'(t)}{\varphi(t)} dt = \log |\varphi(t)| \Big|_a^b, \quad \left(f(x) = \frac{1}{x}, x = \varphi(t) \right).$$

(19.13) Sei $[a, b] \subset]-\frac{\pi}{2}, \frac{\pi}{2}[$. Dann gilt nach (19.12)

$$\int_a^b \tan t dt = \int_a^b \frac{\sin t}{\cos t} dt = -\log \cos t \Big|_a^b.$$

(19.14) Zur Berechnung von $\int_a^b \frac{dx}{1-x^2}$, wobei $-1, 1 \notin [a, b]$, verwendet man die sog. *Partialbruchzerlegung*:

Da $1 - x^2 = (1 - x)(1 + x)$, versucht man $\alpha, \beta \in \mathbb{R}$ so zu bestimmen, dass

$$\frac{1}{1-x^2} = \frac{\alpha}{1-x} + \frac{\beta}{1+x},$$

d.h.

$$\frac{1}{1-x^2} = \frac{(\alpha + \beta) + (\alpha - \beta)x}{1-x^2}.$$

Man erhält $\alpha = \beta = \frac{1}{2}$. Damit folgt

$$\begin{aligned} \int_a^b \frac{dx}{1-x^2} &= \frac{1}{2} \left(\int_a^b \frac{dx}{1-x} + \int_a^b \frac{dx}{1+x} \right) = \frac{1}{2} \left(\int_a^b \frac{dx}{1+x} - \int_a^b \frac{dx}{x-1} \right) \\ &= \frac{1}{2} (\log|x+1| - \log|x-1|) \Big|_a^b = \frac{1}{2} \log \left| \frac{x+1}{x-1} \right| \Big|_a^b. \end{aligned}$$

(19.15) Sei $-1 \leq a < b \leq 1$. Durch die Substitution $x = \sin t$ erhält man mit $u := \arcsin a$, $v := \arcsin b$

$$\int_a^b \sqrt{1-x^2} dx = \int_u^v \sqrt{1-\sin^2 t} d\sin t = \int_u^v \cos^2 t dt.$$

Wegen

$$\cos^2 t = \left(\frac{e^{it} + e^{-it}}{2} \right)^2 = \frac{1}{4} (e^{2it} + e^{-2it}) + \frac{1}{2} = \frac{1}{2} (\cos 2t + 1)$$

folgt weiter

$$\int_a^b \sqrt{1-x^2} dx = \frac{1}{2} \int_u^v (\cos 2t + 1) dt = \frac{1}{4} \sin 2t \Big|_u^v + \frac{1}{2} t \Big|_u^v.$$

Da $\sin 2t = 2 \sin t \cos t = 2 \sin t \sqrt{1-\sin^2 t}$, gilt

$$\sin 2t \Big|_u^v = 2x \sqrt{1-x^2} \Big|_a^b.$$

Also erhält man insgesamt

$$\int_a^b \sqrt{1-x^2} dx = \frac{1}{2} \left(\arcsin x + x \sqrt{1-x^2} \right) \Big|_a^b.$$

Insbesondere für $a = -1$ und $b = 1$ ergibt sich

$$\int_{-1}^1 \sqrt{1-x^2} dx = \frac{1}{2} \arcsin x \Big|_{-1}^{+1} = \frac{\pi}{2},$$

was die Fläche des Halbkreises vom Radius 1 darstellt (Bild 19.1).

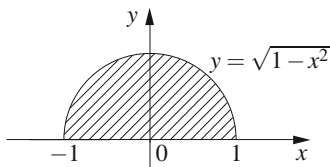


Bild 19.1

(19.16) Zur Berechnung von $\int \frac{dx}{\sqrt{1+x^2}}$ verwenden wir die Substitution $x = \sinh t = \frac{1}{2}(e^t - e^{-t})$. Da

$$d \sinh t = \cosh t \, dt,$$

$$\cosh^2 t - \sinh^2 t = 1, \quad (\text{Aufgabe 10.1}),$$

$$\operatorname{Arsinh} x = \log(x + \sqrt{1+x^2}), \quad (\text{Aufgabe 12.2}),$$

folgt mit $u := \operatorname{Arsinh} a$, $v := \operatorname{Arsinh} b$

$$\begin{aligned} \int_a^b \frac{dx}{\sqrt{1+x^2}} &= \int_u^v \frac{d \sinh t}{\sqrt{1+\sinh^2 t}} = \int_u^v \frac{\cosh t}{\cosh t} dt = t \Big|_u^v \\ &= \log(x + \sqrt{1+x^2}) \Big|_a^b. \end{aligned}$$

(19.17) Berechnung von $\int_a^b \frac{dx}{\sqrt{x^2-1}}$, $(a, b > 1)$.

Wir substituieren $x = \cosh t = \frac{1}{2}(e^t + e^{-t})$. Da

$$d \cosh t = \sinh t \, dt$$

$$\operatorname{Arcosh} x = \log(x + \sqrt{x^2-1}), \quad (\text{Aufgabe 12.2}),$$

folgt mit $u := \operatorname{Arcosh} a$, $v := \operatorname{Arcosh} b$

$$\int_a^b \frac{dx}{\sqrt{x^2-1}} = \int_u^v \frac{\sinh t}{\sinh t} dt = t \Big|_u^v = \log(x + \sqrt{x^2-1}) \Big|_a^b.$$

Partielle Integration

Neben der Substitutionsregel ist die partielle Integration ein weiteres nützliches Hilfsmittel zur Auswertung von Integralen.

Satz 5 (Partielle Integration). *Seien $f, g : [a, b] \rightarrow \mathbb{R}$ zwei stetig differenzierbare Funktionen. Dann gilt*

$$\int_a^b f(x)g'(x) dx = f(x)g(x) \Big|_a^b - \int_a^b g(x)f'(x) dx.$$

Eine Kurzschreibweise für diese Formel ist

$$\int f dg = fg - \int g df.$$

Beweis. Für $F := fg$ gilt nach der Produktregel

$$F'(x) = f'(x)g(x) + f(x)g'(x),$$

also nach Satz 3

$$\int_a^b f'(x)g(x) dx + \int_a^b f(x)g'(x) dx = F(x) \Big|_a^b = f(x)g(x) \Big|_a^b,$$

woraus die Behauptung folgt.

Beispiele

(19.18) Seien $a, b > 0$. Zur Berechnung von $\int_a^b \log x dx$ setzen wir $f(x) = \log x$, $g(x) = x$.

$$\begin{aligned} \int_a^b \log x dx &= x \log x \Big|_a^b - \int_a^b x d \log x = x \log x \Big|_a^b - \int_a^b dx \\ &= x (\log x - 1) \Big|_a^b. \end{aligned}$$

(19.19) Berechnung von $\int \arctan x dx$.

$$\int \arctan x dx = x \arctan x - \int x d \arctan x.$$

Da $\frac{d}{dx} \arctan x = \frac{1}{1+x^2}$, folgt

$$\begin{aligned} \int x d \arctan x &= \int \frac{x}{1+x^2} dx = (\text{Substitution } t = x^2) \\ &= \frac{1}{2} \int \frac{dt}{1+t} = \frac{1}{2} \log(1+t) = \frac{1}{2} \log(1+x^2) \end{aligned}$$

Also gilt

$$\int \arctan x dx = x \arctan x - \frac{1}{2} \log(1+x^2).$$

(19.20) Berechnung von $\int \arcsin x dx$, $(-1 < x < 1)$.

$$\int \arcsin x dx = x \arcsin x - \int x d \arcsin x.$$

Nun ist

$$\begin{aligned} \int x d \arcsin x &= \int \frac{x}{\sqrt{1-x^2}} dx = \quad (t = 1-x^2, dt = -2x dx) \\ &= -\frac{1}{2} \int \frac{dt}{\sqrt{t}} = -\sqrt{t} = -\sqrt{1-x^2}, \end{aligned}$$

also

$$\int \arcsin x dx = x \arcsin x + \sqrt{1-x^2}.$$

Eine zweite Methode zur Berechnung von $\int x d \arcsin x$ liefert die Substitution $t = \arcsin x$:

$$\int x d \arcsin x = \int \sin t dt = -\cos t = -\sqrt{1-\sin^2 t} = -\sqrt{1-x^2}.$$

(Es ist $\cos t \geq 0$, da $-\frac{\pi}{2} \leq t \leq \frac{\pi}{2}$.)

(19.21) Sei $t \neq 0$ ein reeller Parameter. Durch zweimalige Anwendung der partiellen Integration berechnet man

$$\begin{aligned} \int e^{tx} \sin x dx &= \frac{1}{t} e^{tx} \sin x - \frac{1}{t} \int e^{tx} \cos x dx \\ &= \frac{1}{t} e^{tx} \sin x - \frac{1}{t^2} e^{tx} \cos x - \frac{1}{t^2} \int e^{tx} \sin x dx. \end{aligned}$$

Diese Gleichung kann man nach $\int e^{tx} \sin x dx$ auflösen und erhält

$$\int e^{tx} \sin x dx = \frac{e^{tx}}{1+t^2} (t \sin x - \cos x).$$

(19.22) Mithilfe der partiellen Integration kann man manchmal für Integrale, die von einem ganzzahligen Parameter abhängen, Rekursionsformeln herleiten. Als Beispiel betrachten wir für $m \geq 1$ das Integral

$$I_m := \int \frac{dx}{(1+x^2)^m}.$$

Partielle Integration ergibt

$$\begin{aligned} \int \frac{1}{(1+x^2)^m} dx &= \frac{x}{(1+x^2)^m} - \int x d\left(\frac{1}{(1+x^2)^m}\right) \\ &= \frac{x}{(1+x^2)^m} + 2m \int \frac{x^2}{(1+x^2)^{m+1}} dx \\ &= \frac{x}{(1+x^2)^m} + 2m \int \frac{dx}{(1+x^2)^m} - 2m \int \frac{dx}{(1+x^2)^{m+1}}. \end{aligned}$$

Dies bedeutet

$$2mI_{m+1} = (2m-1)I_m + \frac{x}{(1+x^2)^m}.$$

Da $I_1 = \arctan x$, kann man daraus I_m für alle $m \geq 1$ berechnen. Speziell für $m=1$ erhält man aus der obigen Formel

$$\int \frac{dx}{(1+x^2)^2} = \frac{1}{2} \left(\arctan x + \frac{x}{1+x^2} \right),$$

was man auch direkt durch Differenzieren der rechten Seite bestätigen kann.

(19.23) Als weiteres Beispiel für eine Rekursionsformel behandeln wir die Integrale

$$I_m := \int \sin^m x dx.$$

Partielle Integration liefert für $m \geq 2$

$$\begin{aligned} I_m &= - \int \sin^{m-1} x d \cos x \\ &= -\cos x \sin^{m-1} x + (m-1) \int \cos^2 x \sin^{m-2} x dx \\ &= -\cos x \sin^{m-1} x + (m-1) \int (1 - \sin^2 x) \sin^{m-2} x dx \\ &= -\cos x \sin^{m-1} x + (m-1) I_{m-2} - (m-1) I_m. \end{aligned}$$

Diese Gleichung kann man nach I_m auflösen und erhält

$$I_m = -\frac{1}{m} \cos x \sin^{m-1} x + \frac{m-1}{m} I_{m-2}.$$

Da

$$I_0 = \int \sin^0 x dx = x, \quad I_1 = \int \sin x dx = -\cos x,$$

kann man damit rekursiv I_m für alle natürlichen Zahlen m berechnen.

(19.24) Wir wollen das vorangehende Beispiel für das bestimmte Integral

$$A_m := \int_0^{\pi/2} \sin^m x dx$$

ausführen. Es ist $A_0 = \frac{\pi}{2}$, $A_1 = 1$ und

$$A_m = \frac{m-1}{m} A_{m-2} \quad \text{für } m \geq 2.$$

Man erhält

$$A_{2n} = \frac{(2n-1)(2n-3) \cdots 3 \cdot 1}{2n \cdot (2n-2) \cdots 4 \cdot 2} \cdot \frac{\pi}{2},$$

$$A_{2n+1} = \frac{2n \cdot (2n-2) \cdots 4 \cdot 2}{(2n+1) \cdot (2n-1) \cdots 5 \cdot 3}.$$

Folgerung (Wallis'sches Produkt). $\frac{\pi}{2}$ kann durch folgendes unendliche Produkt dargestellt werden:

$$\frac{\pi}{2} = \prod_{n=1}^{\infty} \frac{4n^2}{4n^2 - 1}.$$

Beweis. Wegen $\sin^{2n+2} x \leq \sin^{2n+1} x \leq \sin^{2n} x$ für $x \in [0, \frac{\pi}{2}]$ gilt

$$A_{2n+2} \leq A_{2n+1} \leq A_{2n}.$$

Da $\lim_{n \rightarrow \infty} \frac{A_{2n+2}}{A_{2n}} = \lim_{n \rightarrow \infty} \frac{2n+1}{2n+2} = 1$, gilt auch $\lim_{n \rightarrow \infty} \frac{A_{2n+1}}{A_{2n}} = 1$.

Nun ist

$$\frac{A_{2n+1}}{A_{2n}} = \frac{2n \cdot 2n \cdots 4 \cdot 2 \cdot 2}{(2n+1)(2n-1) \cdots 3 \cdot 3 \cdot 1} \cdot \frac{2}{\pi} = \prod_{k=1}^n \frac{4k^2}{4k^2 - 1} \cdot \frac{2}{\pi}.$$

Grenzübergang $n \rightarrow \infty$ liefert die Behauptung.

Bemerkung. Das Wallis'sche Produkt ist für die praktische Berechnung von π nicht besonders gut geeignet, da es langsam konvergiert. Z.B. ist

$$\prod_{n=1}^{1000} \frac{4n^2}{4n^2 - 1} = 1.57040 \dots,$$

verglichen mit dem exakten Wert von $\frac{\pi}{2} = 1.5707963 \dots$. Die Formel wird uns aber gute Dienste leisten bei der Untersuchung der Gamma-Funktion und beim Beweis der Stirlingschen Formel (§ 20).

Als weitere Anwendung der partiellen Integration beweisen wir:

Satz 6 (Riemannsches Lemma). *Sei $f: [a, b] \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion. Für $k \in \mathbb{R}$ sei*

$$F(k) := \int_a^b f(x) \sin kx dx.$$

Dann gilt $\lim_{|k| \rightarrow \infty} F(k) = 0$.

Beweis. Für $k \neq 0$ ergibt sich durch partielle Integration

$$F(k) = -f(x) \frac{\cos kx}{k} \Big|_a^b + \frac{1}{k} \int_a^b f'(x) \cos kx dx.$$

Da f und f' auf $[a, b]$ stetig sind, gibt es eine Konstante $M \geq 0$, so dass

$$|f(x)| \leq M \quad \text{und} \quad |f'(x)| \leq M \quad \text{für alle } x \in [a, b].$$

Damit ergibt sich die Abschätzung

$$|F(k)| \leq \frac{2M}{|k|} + \frac{M(b-a)}{|k|},$$

woraus die Behauptung folgt.

Bemerkung. Das Riemannsche Lemma gilt auch unter der schwächeren Voraussetzung, dass f nur Riemann-integrierbar ist, siehe Aufgabe 19.9.

(19.25) Als Beispiel für Satz 6 beweisen wir die Formel

$$\sum_{k=1}^{\infty} \frac{\sin kx}{k} = \frac{\pi - x}{2} \quad \text{für } 0 < x < 2\pi.$$

Beweis. Da $\int_{\pi}^x \cos kt dt = \frac{\sin kx}{k}$ und

$$\sum_{k=1}^n \cos kt = \frac{\sin(n + \frac{1}{2})t}{2 \sin \frac{1}{2}t} - \frac{1}{2}, \quad (\text{Hilfssatz aus §18}),$$

folgt

$$\sum_{k=1}^n \frac{\sin kx}{k} = \int_{\pi}^x \frac{\sin(n + \frac{1}{2})t}{2 \sin \frac{1}{2}t} dt - \frac{1}{2}(x - \pi).$$

Nach Satz 6 gilt für

$$F_n(x) := \int_{\pi}^x \frac{1}{2 \sin \frac{1}{2}t} \sin(n + \frac{1}{2})t dt, \quad (0 < x < 2\pi),$$

dass $\lim_{n \rightarrow \infty} F_n(x) = 0$. Daraus folgt die Behauptung.

Spezialfall. Setzt man in der bewiesenen Formel $x = \pi/2$, so erhält man die Leibniz'sche Reihe

$$\frac{\pi}{4} = \sum_{k=0}^{\infty} \frac{(-1)^k}{2k+1} = 1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} + \frac{1}{9} - \frac{1}{11} \pm \dots$$

Satz 7 (Trapez-Regel). *Sei $f: [0, 1] \rightarrow \mathbb{R}$ eine zweimal stetig differenzierbare Funktion. Dann ist*

$$\int_0^1 f(x) dx = \frac{1}{2}(f(0) + f(1)) - R,$$

wobei für das Restglied gilt

$$R = \frac{1}{2} \int_0^1 x(1-x) f''(x) dx = \frac{1}{12} f''(\xi)$$

für ein $\xi \in [0, 1]$.

Beweis. Sei $\varphi(x) := \frac{1}{2}x(1-x)$. Es gilt $\varphi'(x) = \frac{1}{2} - x$ und $\varphi''(x) = -1$. Durch zweimalige partielle Integration erhält man

$$\begin{aligned} R &= \int_0^1 \varphi(x) f''(x) dx = \varphi(x) f'(x) \Big|_0^1 - \int_0^1 \varphi'(x) f'(x) dx \\ &= -\varphi'(x) f(x) \Big|_0^1 + \int_0^1 \varphi''(x) f(x) dx \\ &= \frac{1}{2}(f(0) + f(1)) - \int_0^1 f(x) dx. \end{aligned}$$

Andrerseits kann man wegen $\varphi(x) \geq 0$ für alle $x \in [0, 1]$, auf das Integral für R

den Mittelwertsatz anwenden und erhält ein $\xi \in [0, 1]$ mit

$$R = \int_0^1 \varphi(x) f''(x) dx = f''(\xi) \int_0^1 \varphi(x) dx = \frac{1}{12} f''(\xi), \quad \text{q.e.d.}$$

Bemerkung. Der Name Trapez-Regel kommt daher, dass der Ausdruck $\frac{1}{2}(f(0) + f(1))$ bei positivem f die Fläche des Trapezes mit den Ecken $(0, 0)$, $(1, 0)$, $(0, f(0))$ und $(1, f(1))$ darstellt (Bild 19.2). Man sieht an der Figur auch, warum das Korrekturglied $-\frac{1}{12}f''(\xi)$ mit einem Minuszeichen versehen ist, denn für eine konvexe Funktion (für die $f'' \geq 0$) ist die Fläche des Trapezes größergleich dem Integral.

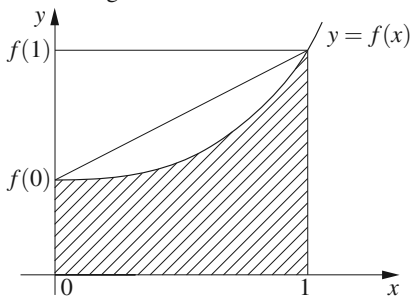


Bild 19.2

Corollar. Es sei $f: [a, b] \rightarrow \mathbb{R}$ eine zweimal stetig differenzierbare Funktion und

$$K := \sup \{ |f''(x)| : x \in [a, b] \}.$$

Sei $n \geq 1$ eine natürliche Zahl und $h := \frac{b-a}{n}$. Dann gilt

$$\int_a^b f(x) dx = \left(\frac{1}{2}f(a) + \sum_{v=1}^{n-1} f(a+vh) + \frac{1}{2}f(b) \right) h + R$$

mit $|R| \leq \frac{K}{12}(b-a)h^2$.

Bemerkung. Lässt man also die Anzahl n der Teilpunkte gegen unendlich gehen, geht der Fehler gegen null, und zwar wird wegen des Gliedes h^2 in der Fehlerabschätzung eine Verdopplung der Anzahl der Teilpunkte zu einer etwa vierfachen Genauigkeit führen. Man kann das Corollar als eine quantitative Präzisierung des Satzes über die Riemannschen Summen (§18, Satz 8) für den Fall zweimal stetig differenzierbarer Funktionen ansehen.

Ist die zu integrierende Funktion 4-mal stetig differenzierbar, so gibt es mit der Simpsonschen Regel (siehe Aufgabe 19.13) eine Näherungsformel für das Integral, bei der das Restglied sogar mit h^4 gegen Null geht.

Beweis. Durch Variablentransformation erhält man aus Satz 7

$$\int_{a+vh}^{a+(v+1)h} f(x) dx = \frac{h}{2} (f(a+vh) + f(a+(v+1)h)) - \frac{h^3}{12} f''(\xi)$$

mit $\xi \in [a+vh, a+(v+1)h]$. Summation über v ergibt die Behauptung.

AUFGABEN

19.1. Seien $a, b \in \mathbb{R}_+^*$. Man berechne den Flächeninhalt der Ellipse

$$E := \left\{ (x, y) \in \mathbb{R}^2 : \frac{x^2}{a^2} + \frac{y^2}{b^2} \leq 1 \right\}.$$

19.2. Für $n, m \in \mathbb{N}$ berechne man die Integrale

$$\int_0^{2\pi} \sin nx \sin mx dx, \quad \int_0^{2\pi} \sin nx \cos mx dx, \quad \int_0^{2\pi} \cos nx \cos mx dx.$$

19.3. Man berechne das Integral

$$\int \sqrt{1+x^2} dx.$$

19.4. Man bestimme eine Rekursionsformel für die Integrale

$$I_m := \int \frac{dx}{\sqrt{1+x^{2m}}}, \quad m \in \mathbb{N}.$$

19.5. Man berechne das Integral

$$\int \frac{dx}{ax^2 + bx + c}$$

in Abhängigkeit von $a, b, c \in \mathbb{R}$, $a \neq 0$.

19.6. Man berechne das Integral $\int \frac{dx}{1+x^4}$.

Anleitung. Man benutze $1+x^4 = (1+\sqrt{2}x+x^2)(1-\sqrt{2}x+x^2)$ und stelle eine Partialbruchzerlegung

$$\frac{1}{1+x^4} = \frac{ax+b}{1+\sqrt{2}x+x^2} + \frac{cx+d}{1-\sqrt{2}x+x^2}$$

her.

19.7. Man berechne die Integrale

$$\int x \sin x dx, \quad \int x^2 \cos x dx, \quad \int x^3 e^x dx.$$

19.8. Für $m \in \mathbb{Z}$ berechne man das Integral

$$\int x^m \log x dx \quad (x > 0).$$

19.9. Man zeige: Das Riemannsche Lemma (Satz 6)

$$\lim_{k \rightarrow \infty} \int_a^b f(x) \sin kx dx = 0$$

gilt auch unter der schwächeren Voraussetzung, dass $f: [a, b] \rightarrow \mathbb{R}$ nur Riemann-integrierbar ist.

Anleitung. Man behandle zunächst den Fall, dass f eine Treppenfunktion ist und führe den allgemeinen Fall durch Approximation darauf zurück.

19.10. Es seien P_n die Legendre-Polynome

$$P_n(x) = \frac{1}{2^n n!} \left(\frac{d}{dx} \right)^n (x^2 - 1)^n,$$

vgl. Aufgabe 16.4. Man beweise mittels partieller Integration

$$\text{i) } \int_{-1}^1 P_n(x) P_m(x) dx = 0 \quad \text{für } n \neq m.$$

$$\text{ii) } \int_{-1}^1 P_n(x)^2 dx = \frac{2}{2n+1}.$$

19.11. Es sei N eine vorgegebene natürliche Zahl. Man beweise:

a) Jedes Polynom f vom Grad $\leq N$ lässt sich als Linearkombination der Legendre-Polynome P_k , $k = 0, 1, \dots, N$, darstellen:

$$f(x) = \sum_{k=0}^N c_k P_k(x), \quad \text{wobei} \quad c_k = \frac{2n+1}{2} \int_{-1}^1 f(x) P_k(x) dx.$$

b) Für jedes Polynom g vom Grad $< N$ gilt

$$\int_{-1}^1 g(x) P_N(x) dx = 0.$$

c) Seien $x_1, x_2, \dots, x_N$ die Nullstellen des Polynoms P_N (vgl. Aufgabe 16.4) und sei f ein Polynom vom Grad $\leq 2N - 1$ mit

$$f(x_k) = 0 \quad \text{für } k = 1, 2, \dots, N.$$

Dann gilt

$$\int_{-1}^1 f(x) dx = 0.$$

d) Für $n = 1, 2, \dots, N$ sei

$$\gamma_n := \int_{-1}^1 L_n(x) dx, \quad \text{wobei} \quad L_n(x) := \prod_{\substack{k=1 \\ k \neq n}}^N \frac{x - x_k}{x_n - x_k}.$$

Dann gilt für jedes Polynom f vom Grad $\leq 2N - 1$

$$\int_{-1}^1 f(x) dx = \sum_{k=1}^N \gamma_k f(x_k) \quad (\text{Gauß'sche Quadratur-Formel}).$$

e) Man berechne die x_k und γ_k für die Fälle $N = 1, 2, 3$.

19.12. a) Sei $f : [-\frac{1}{2}, \frac{1}{2}] \rightarrow \mathbb{R}$ eine zweimal stetig differenzierbare Funktion. Man zeige: Es gibt ein $\xi \in [-\frac{1}{2}, \frac{1}{2}]$, so dass

$$\begin{aligned} \int_{-1/2}^{1/2} f(x) dx &= f(0) + \frac{1}{2} \int_{-1/2}^{1/2} (|x| - \frac{1}{2})^2 f''(x) dx \\ &= f(0) + \frac{1}{24} f''(\xi). \end{aligned}$$

b) Sei $f: [a, b] \rightarrow \mathbb{R}$ eine zweimal stetig differenzierbare Funktion und

$$K_2 := \sup\{|f''(x)| : x \in [a, b]\}.$$

Weiter sei $n > 0$ eine natürliche Zahl und $h := (b-a)/n$. Man zeige

$$\int_a^b f(x) dx = h \sum_{v=0}^{n-1} f\left(a + \left(v + \frac{1}{2}\right)h\right) + R$$

mit $|R| \leq \frac{K_2}{24}(b-a)h^2$.

19.13. Sei $\psi: \mathbb{R} \rightarrow \mathbb{R}$ die wie folgt definierte Funktion:

$$\psi(x) := \frac{1}{18}(|x| - 1)^3 + \frac{1}{24}(|x| - 1)^4.$$

a) Man zeige für jede 4-mal stetig differenzierbare Funktion $f: [-1, 1] \rightarrow \mathbb{R}$

$$\int_{-1}^1 f(x) dx = \frac{1}{3}(f(-1) + 4f(0) + f(1)) + R$$

(*Keplersche Fassregel*), wobei

$$R = \int_{-1}^1 f^{(4)}(x) \psi(x) dx = -\frac{1}{90} f^{(4)}(\xi) \quad \text{für ein } \xi \in [-1, 1].$$

b) Sei $f: [a, b] \rightarrow \mathbb{R}$ eine 4-mal stetig differenzierbare Funktion und

$$K_4 := \sup\{|f^{(4)}(x)| : x \in [a, b]\}.$$

Weiter sei $n > 0$ eine natürliche Zahl und

$$h := (b-a)/2n, \quad x_v := a + vh, \quad y_v := f(x_v).$$

Man beweise die *Simpsonsche Regel*

$$\begin{aligned} \int_a^b f(x) dx &= \frac{h}{3}(y_0 + 4y_1 + 2y_2 + 4y_3 + \dots + 2y_{2n-2} + 4y_{2n-1} + y_{2n}) + R \\ &= \frac{h}{3} \left(f(x_0) + 4 \sum_{v=0}^{n-1} f(x_{2v+1}) + 2 \sum_{v=1}^{n-1} f(x_{2v}) + f(x_{2n}) \right) + R \end{aligned}$$

mit $|R| \leq \frac{K_4}{180}(b-a)h^4$.

§ 20 Uneigentliche Integrale. Die Gamma-Funktion

Der bisher behandelte Integralbegriff ist für manche Anwendungen zu eng. So konnten wir bisher nur über endliche Intervalle integrieren und die Riemann-integrierbaren Funktionen waren notwendig beschränkt. Ist das Integrationsintervall unendlich oder die zu integrierende Funktion nicht beschränkt, so kommt man zu den uneigentlichen Integralen, die unter gewissen Bedingungen als Grenzwerte Riemannscher Integrale definiert werden können. Als Anwendung behandeln wir die Gamma-Funktion, die durch ein uneigentliches Integral definiert ist und die die Fakultät interpoliert.

Uneigentliche Integrale

Wir betrachten drei Fälle.

Fall 1. Eine Integrationsgrenze ist unendlich.

Definition. Sei $f: [a, \infty[\rightarrow \mathbb{R}$ eine Funktion, die über jedem Intervall $[a, R]$, $a < R < \infty$, Riemann-integrierbar ist. Falls der Grenzwert

$$\lim_{R \rightarrow \infty} \int_a^R f(x) dx$$

existiert, heißt das Integral $\int_a^\infty f(x) dx$ konvergent und man setzt

$$\int_a^\infty f(x) dx := \lim_{R \rightarrow \infty} \int_a^R f(x) dx.$$

Analog definiert man das Integral $\int_{-\infty}^a f(x) dx$ für eine Funktion $f:]-\infty, a] \rightarrow \mathbb{R}$.

(20.1) Beispiel. Das Integral $\int_1^\infty \frac{dx}{x^s}$ konvergiert für $s > 1$. Es gilt nämlich

$$\int_1^R \frac{dx}{x^s} = \frac{1}{1-s} \cdot \frac{1}{x^{s-1}} \Big|_1^R = \frac{1}{s-1} \left(1 - \frac{1}{R^{s-1}} \right).$$

Da $\lim_{R \rightarrow \infty} \frac{1}{R^{s-1}} = 0$, folgt

$$\int_1^{\infty} \frac{dx}{x^s} = \frac{1}{s-1} \quad \text{für } s > 1.$$

Andererseits zeigt man: $\int_1^{\infty} \frac{dx}{x^s}$ konvergiert nicht für $s \leq 1$.

Z.B. für $s = 1$ ist $\int_1^R \frac{dx}{x} = \log R$, was für $R \rightarrow \infty$ gegen ∞ strebt.

Fall 2. Der Integrand ist an einer Integrationsgrenze nicht definiert.

Definition. Sei $f:]a, b] \rightarrow \mathbb{R}$ eine Funktion, die über jedem Teilintervall $[a + \varepsilon, b]$, $0 < \varepsilon < b - a$, Riemann-integrierbar ist. Falls der Grenzwert

$$\lim_{\varepsilon \searrow 0} \int_{a+\varepsilon}^b f(x) dx$$

existiert, heißt das Integral $\int_a^b f(x) dx$ konvergent und man setzt

$$\int_a^b f(x) dx := \lim_{\varepsilon \searrow 0} \int_{a+\varepsilon}^b f(x) dx.$$

(20.2) Beispiel. Das Integral $\int_0^1 \frac{dx}{x^s}$ konvergiert für $s < 1$. Es gilt nämlich

$$\int_{\varepsilon}^1 \frac{dx}{x^s} = \frac{1}{1-s} \cdot \frac{1}{x^{s-1}} \Big|_{\varepsilon}^1 = \frac{1}{1-s} (1 - \varepsilon^{1-s}).$$

Da $\lim_{\varepsilon \searrow 0} \varepsilon^{1-s} = 0$, folgt

$$\int_0^1 \frac{dx}{x^s} = \frac{1}{1-s} \quad \text{für } s < 1.$$

Andererseits zeigt man

$$\int_0^1 \frac{dx}{x^s} \quad \text{konvergiert nicht für } s \geq 1.$$

Fall 3. Beide Integrationsgrenzen sind kritisch.

Definition. Sei $f:]a, b[\rightarrow \mathbb{R}$, $a \in \mathbb{R} \cup \{-\infty\}$, $b \in \mathbb{R} \cup \{\infty\}$, eine Funktion, die über jedem kompakten Teilintervall $[\alpha, \beta] \subset]a, b[$ Riemann-integrierbar ist und sei $c \in]a, b[$ beliebig. Falls die beiden uneigentlichen Integrale

$$\int_a^c f(x) dx = \lim_{\alpha \searrow a} \int_{\alpha}^c f(x) dx$$

und

$$\int_c^b f(x) dx = \lim_{\beta \nearrow b} \int_c^{\beta} f(x) dx$$

konvergieren, heißt das Integral $\int_a^b f(x) dx$ konvergent und man setzt

$$\int_a^b f(x) dx = \int_a^c f(x) dx + \int_c^b f(x) dx.$$

Bemerkung. Diese Definition ist unabhängig von der Auswahl von $c \in]a, b[$.

Beispiele

(20.3) Nach (20.1) und (20.2) divergiert das Integral $\int_0^{\infty} \frac{dx}{x^s}$ für jedes $s \in \mathbb{R}$.

(20.4) Das Integral $\int_{-1}^1 \frac{dx}{\sqrt{1-x^2}}$ konvergiert:

$$\begin{aligned} \int_{-1}^1 \frac{dx}{\sqrt{1-x^2}} &= \lim_{\varepsilon \searrow 0} \int_{-1+\varepsilon}^0 \frac{dx}{\sqrt{1-x^2}} + \lim_{\varepsilon \searrow 0} \int_0^{1-\varepsilon} \frac{dx}{\sqrt{1-x^2}} \\ &= -\lim_{\varepsilon \searrow 0} \arcsin(-1+\varepsilon) + \lim_{\varepsilon \searrow 0} \arcsin(1-\varepsilon) \\ &= -\left(-\frac{\pi}{2}\right) + \frac{\pi}{2} = \pi. \end{aligned}$$

(20.5) Das Integral $\int_{-\infty}^{\infty} \frac{dx}{1+x^2}$ konvergiert ebenfalls:

$$\begin{aligned} \int_{-\infty}^{\infty} \frac{dx}{1+x^2} &= \lim_{R \rightarrow \infty} \int_{-R}^0 \frac{dx}{1+x^2} + \lim_{R \rightarrow \infty} \int_0^R \frac{dx}{1+x^2} \\ &= -\lim_{R \rightarrow \infty} \arctan(-R) + \lim_{R \rightarrow \infty} \arctan(R) \\ &= -\left(-\frac{\pi}{2}\right) + \frac{\pi}{2} = \pi. \end{aligned}$$

(20.6) Wir behandeln jetzt noch ein nicht-triviales Beispiel eines uneigentlichen Integrals und beweisen die Dirichletsche Formel

$$\int_0^{\infty} \frac{\sin x}{x} dx = \frac{\pi}{2}.$$

Die Integrationsgrenze 0 ist nicht kritisch, da sich der Intergrand $\frac{\sin x}{x}$ durch 1 stetig in die Stelle $x = 0$ fortsetzen lässt (§14, Corollar zu Satz 5). Um die Konvergenz des Integrals an der oberen Grenze ∞ zu untersuchen, betrachten wir das unbestimmte Integral

$$\text{Si}(x) := \int_0^x \frac{\sin t}{t} dt, \quad 0 \leq x < \infty.$$

Diese Funktion heißt *Integral-Sinus* (oder Sinus integralis) und lässt sich nicht wie im vorigen Beispiel (20.5) durch elementare Funktionen ausdrücken. Da $\frac{\sin x}{x}$ im Intervall $n\pi < x < (n+1)\pi$ für gerades n positiv und für ungerades n negativ ist, hat $\text{Si}(x)$ an den Stellen $x = n\pi$ lokale Maxima bzw. Minima, je nachdem n ungerade oder gerade ist. Setzt man

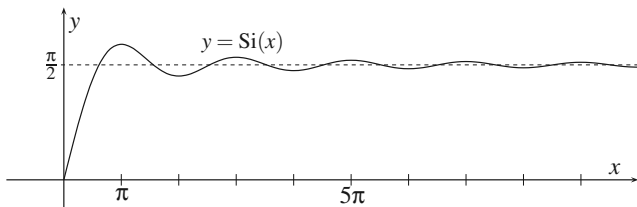
$$a_n := \left| \int_{n\pi}^{(n+1)\pi} \frac{\sin x}{x} dx \right|,$$

so ist die Folge (a_n) monoton fallend mit Limes 0 und es gilt

$$\text{Si}(n\pi) = \sum_{k=0}^n (-1)^k a_k.$$

Die Existenz von $\int_0^{\infty} \frac{\sin x}{x} dx = \lim_{R \rightarrow \infty} \text{Si}(R) = \lim_{n \rightarrow \infty} \text{Si}(n\pi)$ folgt nun aus dem Leibniz'schen Konvergenz-Kriterium für alternierende Reihen.

Bild 20.1 zeigt den Graphen des Integral-Sinus (die x - und y -Achse haben verschiedenen Maßstab!).

**Bild 20.1** Die Funktion Sinus integralis

Wir kommen jetzt zur Berechnung des Limes. Wir gehen in drei Schritten vor.

i) Zunächst folgt für jede positive reelle Zahl λ durch einfache Variablen-Substitution

$$\text{Si}(\lambda\pi/2) = \int_0^{\pi/2} \frac{\sin \lambda x}{x} dx.$$

ii) Sei $g : [0, \pi/2] \rightarrow \mathbb{R}$ die wie folgt definierte Funktion

$$g(x) := \frac{1}{x} - \frac{1}{\sin x} \quad \text{für } x \neq 0, \quad g(0) := 0.$$

Nach (16.4) ist g im Nullpunkt stetig. Aus dem Riemannschen Lemma (§19, Satz 6) folgt

$$\lim_{\lambda \rightarrow \infty} \int_0^{\pi/2} \sin(\lambda x) g(x) dx = 0,$$

also

$$\lim_{\lambda \rightarrow \infty} \int_0^{\pi/2} \frac{\sin \lambda x}{x} dx = \lim_{\lambda \rightarrow \infty} \int_0^{\pi/2} \frac{\sin \lambda x}{\sin x} dx.$$

iii) Für jede positive ganze Zahl n gilt

$$\frac{\sin(2n+1)x}{\sin x} = 1 + 2 \sum_{k=1}^n \cos 2kx,$$

vgl. den Hilfssatz in §18, Seite 214. Daraus folgt

$$\int_0^{\pi/2} \frac{\sin(2n+1)x}{\sin x} dx = \int_0^{\pi/2} 1 \cdot dx = \frac{\pi}{2}.$$

Zusammenfassend ergibt sich

$$\int_0^{\infty} \frac{\sin x}{x} dx = \lim_{n \rightarrow \infty} \int_0^{\pi/2} \frac{\sin(2n+1)x}{x} dx = \lim_{n \rightarrow \infty} \int_0^{\pi/2} \frac{\sin(2n+1)x}{\sin x} dx = \frac{\pi}{2}, \text{ q.e.d.}$$

Integral-Vergleichskriterium für Reihen

Mithilfe der uneigentlichen Integrale kann man manchmal einfach entscheiden, ob eine unendliche Reihe konvergiert oder divergiert.

Satz 1. Sei $f: [1, \infty[\rightarrow \mathbb{R}_+$ eine monoton fallende Funktion. Dann gilt:

$$\sum_{n=1}^{\infty} f(n) \text{ konvergiert} \iff \int_1^{\infty} f(x) dx \text{ konvergiert.}$$

Beweis. Wir definieren Treppenfunktionen $\varphi, \psi: [1, \infty[\rightarrow \mathbb{R}$ durch

$$\left. \begin{aligned} \psi(x) &:= f(n) \\ \varphi(x) &:= f(n+1) \end{aligned} \right\} \quad \text{für } n \leq x < n+1.$$

Da f monoton fallend ist, gilt $\varphi \leq f \leq \psi$, siehe Bild 20.2.

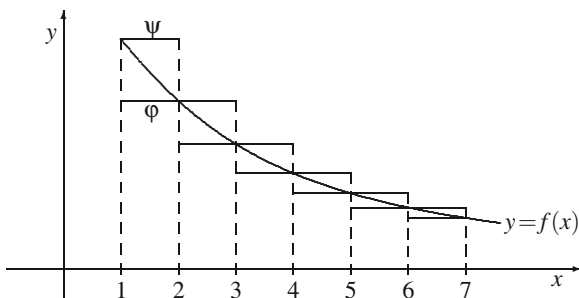


Bild 20.2 Zum Integral-Vergleichskriterium

Integration über das Intervall $[1, N]$ ergibt

$$\sum_{n=2}^N f(n) = \int_1^N \varphi(x) dx \leq \int_1^N f(x) dx \leq \int_1^N \psi(x) dx = \sum_{n=1}^{N-1} f(n).$$

Falls $\int_1^{\infty} f(x) dx$ konvergiert, ist deshalb die Reihe $\sum_{n=1}^{\infty} f(n)$ beschränkt, also konvergent. Falls umgekehrt $\sum_{n=1}^{\infty} f(n)$ als konvergent vorausgesetzt wird, so folgt, dass $\int_1^R f(x) dx$ für $R \rightarrow \infty$ monoton wachsend und beschränkt ist, also konvergiert.

(20.7) Beispiel. Aus (20.1) folgt:

Die Reihe $\sum_{n=1}^{\infty} \frac{1}{n^s}$ konvergiert für $s > 1$ und divergiert für $s \leq 1$.

(Diese Reihe hatten wir schon in (7.2) behandelt.)

Bemerkung. Betrachtet man die Summe der Reihe als Funktion von s , so erhält man die *Riemannsche Zetafunktion*

$$\zeta(s) := \sum_{n=1}^{\infty} \frac{1}{n^s}, \quad (s > 1).$$

(Die wahre Bedeutung dieser Funktion wird erst in der sogenannten Funktionentheorie sichtbar, wo diese Funktion ins Komplexe fortgesetzt wird.)

(20.8) Euler-Mascheronische Konstante. Insbesondere erhält man aus dem vorigen Beispiel für $s = 1$ wieder die Divergenz der harmonischen Reihe $\sum_{n=1}^{\infty} \frac{1}{n}$. Mit dem Integral-Vergleichskriterium kann man auch eine Aussage über das Wachstum der Partialsummen machen. Da $\int_1^N \frac{dx}{x} = \log N$, wächst $\sum_{n=1}^N \frac{1}{n}$ ungefähr so schnell (langsam) wie $\log N$ gegen ∞ . Genauer gilt:

Es gibt eine Konstante $\gamma \in [0, 1]$, so dass

$$\gamma = \lim_{N \rightarrow \infty} \left(\sum_{n=1}^N \frac{1}{n} - \log N \right).$$

Beweis. Mit dem Integral-Vergleichskriterium ergibt sich für $N > 1$

$$\sum_{n=2}^N \frac{1}{n} \leq \int_1^N \frac{dx}{x} = \log N \leq \sum_{n=1}^{N-1} \frac{1}{n}.$$

Daraus folgt

$$\frac{1}{N} \leq \gamma_N := \sum_{n=1}^N \frac{1}{n} - \log N \leq 1.$$

Die Differenz zweier aufeinander folgender Glieder der Folge (γ_N) ist

$$\gamma_{N-1} - \gamma_N = \int_{N-1}^N \frac{dx}{x} - \frac{1}{N} = \int_{N-1}^N \left(\frac{1}{x} - \frac{1}{N} \right) dx > 0.$$

Die Folge der γ_N ist also monoton fallend und durch 0 nach unten beschränkt. Also existiert der Limes

$$\gamma = \lim_{N \rightarrow \infty} \gamma_N = \lim_{N \rightarrow \infty} \left(\sum_{n=1}^N \frac{1}{n} - \log N \right), \quad \text{q.e.d.}$$

Bemerkung. Diese Beziehung lässt sich auch wie folgt schreiben:

$$\sum_{n=1}^N \frac{1}{n} = \log N + \gamma + o(1) \quad \text{für } N \rightarrow \infty.$$

Die Zahl γ heißt die Euler-Mascheronische Konstante; ihr numerischer Wert ist auf 45 Dezimalstellen genau (siehe dazu Aufgabe 23.9)

$$\gamma = 0.57721\,56649\,01532\,86060\,65120\,90082\,40243\,10421\,59335\ldots$$

Es ist unbekannt, ob γ rational, irrational oder (vermutlich) sogar transzendent ist.

(20.9) Alternierende harmonische Reihe

Mithilfe der Euler-Mascheronischen Konstante lässt sich die Summe der alternierenden harmonischen Reihe $\sum_{n=1}^{\infty} (-1)^{n-1} \frac{1}{n}$ wie folgt bestimmen: Es ist

$$\begin{aligned} \sum_{n=1}^{2N} (-1)^{n-1} \frac{1}{n} &= \sum_{n=1}^{2N} \frac{1}{n} - 2 \sum_{n=1}^N \frac{1}{2n} = \sum_{n=1}^{2N} \frac{1}{n} - \sum_{n=1}^N \frac{1}{n} \\ &= \log(2N) + \gamma + o(1) - (\log N + \gamma + o(1)) \\ &= \log 2 + o(1), \end{aligned}$$

$$\text{also} \quad \sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{n} = \log 2.$$

Die Gamma-Funktion

Definition (Eulersche Integraldarstellung der Gamma-Funktion). Für $x > 0$ setzt man

$$\Gamma(x) := \int_0^{\infty} t^{x-1} e^{-t} dt.$$

Bemerkung. Dass dieses uneigentliche Integral konvergiert, folgt nach (20.1) und (20.2) daraus, dass

$$\text{a) } t^{x-1} e^{-t} \leq \frac{1}{t^{1-x}} \quad \text{für alle } t > 0,$$

$$\text{b) } t^{x-1} e^{-t} \leq \frac{1}{t^2} \quad \text{für } t \geq t_0,$$

$$\text{da } \lim_{t \rightarrow \infty} t^{x+1} e^{-t} = 0, \text{ vgl. (12.2).}$$

Satz 2 (Funktionalgleichung). *Es gilt $\Gamma(n+1) = n!$ für alle $n \in \mathbb{N}$ und*

$$x\Gamma(x) = \Gamma(x+1) \quad \text{für alle } x \in \mathbb{R}_+^*.$$

Beweis. Partielle Integration liefert

$$\int_{\varepsilon}^R t^x e^{-t} dt = -t^x e^{-t} \Big|_{t=\varepsilon}^{t=R} + x \int_{\varepsilon}^R t^{x-1} e^{-t} dt.$$

Durch Grenzübergang $\varepsilon \searrow 0$ und $R \rightarrow \infty$ erhält man $\Gamma(x+1) = x\Gamma(x)$. Da

$$\Gamma(1) = \lim_{R \rightarrow \infty} \int_0^R e^{-t} dt = \lim_{R \rightarrow \infty} (1 - e^{-R}) = 1,$$

folgt aus dieser Funktionalgleichung

$$\Gamma(n+1) = n\Gamma(n) = n(n-1)\Gamma(n-1) = n(n-1) \cdot \dots \cdot 1 \cdot \Gamma(1) = n!$$

Bemerkung. Die Funktion $\Gamma: \mathbb{R}_+^* \rightarrow \mathbb{R}$ interpoliert also die Fakultät, die nur für natürliche Zahlen definiert ist. (Dass die Gamma-Funktion so definiert ist, dass nicht $\Gamma(n)$, sondern $\Gamma(n+1)$ gleich $n!$ ist, hat historische Gründe.) Durch diese Eigenschaft und die Funktionalgleichung ist die Gammafunktion aber noch nicht eindeutig bestimmt. Wir brauchen noch eine weitere Eigenschaft, die logarithmische Konvexität, um die Gammafunktion zu charakterisieren.

Definition. Sei $I \subset \mathbb{R}$ ein Intervall. Eine positive Funktion $F: I \rightarrow \mathbb{R}_+^*$ heißt *logarithmisch konvex*, wenn die Funktion $\log F: I \rightarrow \mathbb{R}$ konvex ist.

Übersetzt man die Konvexitätsbedingung für die Funktion $\log F$ mithilfe der Exponentialfunktion auf die Funktion F , so erhält man: F ist genau dann logarithmisch konvex, wenn für alle $x, y \in I$ und $0 < \lambda < 1$ gilt

$$F(\lambda x + (1-\lambda)y) \leq F(x)^\lambda F(y)^{1-\lambda}.$$

Satz 3. *Die Funktion $\Gamma: \mathbb{R}_+^* \rightarrow \mathbb{R}$ ist logarithmisch konvex.*

Beweis. Aus der Integraldarstellung folgt unmittelbar $\Gamma(x) > 0$ für alle $x > 0$.

Seien nun $x, y \in \mathbb{R}_+^*$ und $0 < \lambda < 1$. Wir setzen $p := \frac{1}{\lambda}$ und $q := \frac{1}{1-\lambda}$. Dann gilt $\frac{1}{p} + \frac{1}{q} = 1$. Wir wenden nun auf die Funktionen

$$f(t) := t^{(x-1)/p} e^{-t/p}, \quad g(t) := t^{(y-1)/q} e^{-t/q}$$

die Höldersche Ungleichung (18.6) an:

$$\int_{\varepsilon}^R f(t)g(t) dt \leq \left(\int_{\varepsilon}^R f(t)^p dt \right)^{1/p} \left(\int_{\varepsilon}^R g(t)^q dt \right)^{1/q}.$$

Nun ist

$$\begin{aligned} f(t)g(t) &= t^{\frac{x}{p} + \frac{y}{q} - 1} e^{-t}, \\ f(t)^p &= t^{x-1} e^{-t}, \quad g(t)^q = t^{y-1} e^{-t}. \end{aligned}$$

Damit ergibt die Höldersche Ungleichung nach Grenzübergang $\varepsilon \searrow 0$ und $R \rightarrow \infty$

$$\Gamma\left(\frac{x}{p} + \frac{y}{q}\right) \leq \Gamma(x)^{1/p} \Gamma(y)^{1/q}.$$

Dies zeigt, dass Γ logarithmisch konvex ist.

Bemerkung. Da jede auf einem offenen Intervall konvexe Funktion stetig ist (vgl. Aufgabe 16.5), ergibt sich aus Satz 3 insbesondere, dass die Gamma-Funktion auf ganz $\mathbb{R}_+^*$ stetig ist.

Satz 4 (H. Bohr / J. Møllerup). Sei $F: \mathbb{R}_+^* \rightarrow \mathbb{R}_+^*$ eine Funktion mit folgenden Eigenschaften:

- a) $F(1) = 1$,
- b) $F(x+1) = xF(x)$ für alle $x \in \mathbb{R}_+^*$,
- c) F ist logarithmisch konvex.

Dann gilt $F(x) = \Gamma(x)$ für alle $x \in \mathbb{R}_+^*$.

Beweis. Da die Γ -Funktion die Eigenschaften a) bis c) hat, genügt es zu zeigen, dass eine Funktion F mit a) bis c) eindeutig bestimmt ist.

Aus der Funktionalgleichung b) folgt

$$F(x+n) = F(x)x(x+1) \cdots (x+n-1)$$

für alle $x > 0$ und alle natürlichen Zahlen $n \geq 1$. Insbesondere folgt daraus $F(n+1) = n!$ für alle $n \in \mathbb{N}$. Es genügt daher zu beweisen, dass $F(x)$ für $0 < x < 1$ eindeutig bestimmt ist. Wegen

$$n+x = (1-x)n + x(n+1)$$

folgt aus der logarithmischen Konvexität

$$F(n+x) \leq F(n)^{1-x} F(n+1)^x = F(n)^{1-x} F(n)^x n^x = (n-1)! n^x.$$

Aus $n+1 = x(n+x) + (1-x)(n+1+x)$ folgt ebenso

$$n! = F(n+1) \leq F(n+x)^x F(n+1+x)^{1-x} = F(n+x)(n+x)^{1-x}.$$

Kombiniert man beide Ungleichungen, erhält man

$$n!(n+x)^{x-1} \leq F(n+x) \leq (n-1)!n^x$$

und weiter

$$\begin{aligned} a_n(x) &:= \frac{n!(n+x)^{x-1}}{x(x+1) \cdots (x+n-1)} \leq F(x) \\ &\leq \frac{(n-1)!n^x}{x(x+1) \cdots (x+n-1)} =: b_n(x). \end{aligned}$$

Da $\frac{b_n(x)}{a_n(x)} = \frac{(n+x)n^x}{n(n+x)^x}$ für $n \rightarrow \infty$ gegen 1 konvergiert, folgt

$$F(x) = \lim_{n \rightarrow \infty} \frac{(n-1)!n^x}{x(x+1) \cdots (x+n-1)},$$

F ist also eindeutig bestimmt.

Satz 5 (Gauß'sche Limesdarstellung der Gamma-Funktion). *Für alle $x > 0$ gilt*

$$\Gamma(x) = \lim_{n \rightarrow \infty} \frac{n!n^x}{x(x+1) \cdots (x+n)}.$$

Beweis. Da $\lim_{n \rightarrow \infty} \frac{n}{x+n} = 1$, folgt die behauptete Gleichung für $0 < x < 1$ aus der im vorangehenden Beweis hergeleiteten Beziehung

$$\Gamma(x) = \lim_{n \rightarrow \infty} \frac{(n-1)!n^x}{x(x+1) \cdots (x+n-1)}.$$

Sie ist außerdem trivialerweise für $x = 1$ richtig. Es genügt also zu zeigen: Gilt die Formel für ein x , so auch für $y := x+1$. Nun ist

$$\begin{aligned} \Gamma(y) &= \Gamma(x+1) = x\Gamma(x) = \lim_{n \rightarrow \infty} \frac{n!n^x}{(x+1) \cdots (x+n)} \\ &= \lim_{n \rightarrow \infty} \frac{n!n^{y-1}}{y(y+1) \cdots (y+n-1)} \\ &= \lim_{n \rightarrow \infty} \frac{n!n^y}{n(y(y+1) \cdots (y+n-1)(y+n))}. \end{aligned}$$

Damit ist Satz 5 bewiesen.

Corollar (Weierstraß'sche Produktdarstellung der Gamma-Funktion).

Für alle $x > 0$ gilt

$$\frac{1}{\Gamma(x)} = x e^{\gamma x} \prod_{n=1}^{\infty} \left(1 + \frac{x}{n}\right) e^{-x/n},$$

wobei γ die Euler-Mascheronische Konstante ist.

Beweis. Aus Satz 5 folgt

$$\begin{aligned} \frac{1}{\Gamma(x)} &= x \lim_{N \rightarrow \infty} \frac{(x+1) \cdot \dots \cdot (x+N)}{N!} N^{-x} \\ &= x \lim_{N \rightarrow \infty} \left(1 + \frac{x}{1}\right) \cdot \dots \cdot \left(1 + \frac{x}{N}\right) \exp(-x \log N) \\ &= x \lim_{N \rightarrow \infty} \left(\prod_{n=1}^N \left(1 + \frac{x}{n}\right) e^{-x/n} \right) \exp\left(\sum_{n=1}^N \frac{x}{n} - x \log N\right). \end{aligned}$$

Da $\lim_{N \rightarrow \infty} \left(\sum_{n=1}^N \frac{x}{n} - x \log N\right) = x\gamma$, vgl. (20.8), konvergiert auch das unendliche Produkt und man erhält

$$\frac{1}{\Gamma(x)} = x e^{\gamma x} \prod_{n=1}^{\infty} \left(1 + \frac{x}{n}\right) e^{-x/n}, \quad \text{q.e.d.}$$

(20.10) Wir zeigen als Anwendung von Satz 5, dass $\Gamma(\frac{1}{2}) = \sqrt{\pi}$.

Beweis. Wir können $\Gamma(\frac{1}{2})$ auf zwei Weisen darstellen:

$$\begin{aligned} \Gamma(\tfrac{1}{2}) &= \lim_{n \rightarrow \infty} \frac{n! \sqrt{n}}{\frac{1}{2}(1 + \frac{1}{2})(2 + \frac{1}{2}) \cdot \dots \cdot (n + \frac{1}{2})}, \\ \Gamma(\tfrac{1}{2}) &= \lim_{n \rightarrow \infty} \frac{n! \sqrt{n}}{(1 - \frac{1}{2})(2 - \frac{1}{2}) \cdot \dots \cdot (n - \frac{1}{2})(n + \frac{1}{2})}. \end{aligned}$$

Multiplikation ergibt

$$\begin{aligned} \Gamma(\tfrac{1}{2})^2 &= \lim_{n \rightarrow \infty} \frac{2n}{n + \frac{1}{2}} \cdot \frac{(n!)^2}{(1 - \frac{1}{4})(4 - \frac{1}{4}) \cdot \dots \cdot (n^2 - \frac{1}{4})} \\ &= 2 \lim_{n \rightarrow \infty} \prod_{k=1}^n \frac{k^2}{k^2 - \frac{1}{4}} = \pi, \end{aligned}$$

wobei das Wallis'sche Produkt (19.24) benutzt wurde. Also ist $\Gamma(\frac{1}{2}) = \sqrt{\pi}$.

Daraus kann man $\Gamma(n + \frac{1}{2})$ für alle natürlichen Zahlen n berechnen, denn aus der Funktionalgleichung folgt $(n + \frac{1}{2})\Gamma(n + \frac{1}{2}) = \Gamma(n + 1 + \frac{1}{2})$. Damit erhalten wir folgende Wertetabelle für die Gamma-Funktion:

x	$\Gamma(x)$
0.5	$\sqrt{\pi} = 1.77245\dots$
1	$0! = 1$
1.5	$\frac{1}{2}\sqrt{\pi} = 0.88622\dots$
2	$1! = 1$
2.5	$\frac{3}{4}\sqrt{\pi} = 1.32934\dots$
3	$2! = 2$
3.5	$\frac{15}{8}\sqrt{\pi} = 3.32335\dots$
4	$3! = 6$

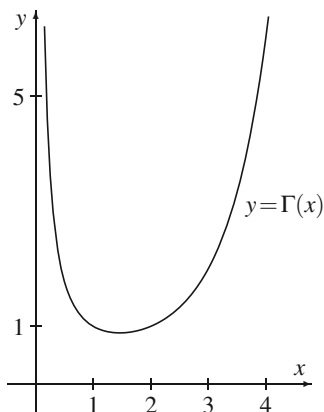


Bild 20.3 Gamma-Funktion

Für $x \searrow 0$ strebt $\Gamma(x)$ gegen unendlich, da

$$\lim_{x \searrow 0} x\Gamma(x) = \lim_{x \searrow 0} \Gamma(x+1) = \Gamma(1) = 1.$$

Dies bedeutet, dass sich $\Gamma(x)$ asymptotisch für $x \searrow 0$ wie die Funktion $x \mapsto 1/x$ verhält.

(20.11) Mithilfe des Wertes von $\Gamma(\frac{1}{2})$ können wir das folgende uneigentliche Integral berechnen:

$$\int_{-\infty}^{\infty} e^{-x^2} dx = \sqrt{\pi}.$$

Beweis. Die Substitution $x = t^{1/2}$, $dx = \frac{1}{2}t^{-1/2}dt$ liefert

$$\int_{\varepsilon}^R e^{-x^2} dx = \frac{1}{2} \int_{\varepsilon^2}^{R^2} t^{-1/2} e^{-t} dt,$$

also ergibt sich durch Grenzübergang $\varepsilon \searrow 0$, $R \rightarrow \infty$

$$\int_0^{\infty} e^{-x^2} dx = \frac{1}{2} \int_0^{\infty} t^{-1/2} e^{-t} dt = \frac{1}{2} \Gamma\left(\frac{1}{2}\right) = \frac{1}{2} \sqrt{\pi}, \quad \text{q.e.d.}$$

Stirlingsche Formel

Wir leiten jetzt noch eine nützliche Formel für das asymptotische Verhalten von $n!$ für große n her. Dabei nennt man zwei Folgen $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ nichtverschwindender Zahlen *asymptotisch gleich*, in Zeichen $a_n \sim b_n$, falls

$$\lim_{n \rightarrow \infty} \frac{a_n}{b_n} = 1.$$

Man beachte, dass nicht vorausgesetzt wird, dass die beiden Folgen (a_n) und (b_n) konvergieren und dass auch im Allgemeinen die Folge der Differenzen $(a_n - b_n)$ nicht konvergiert.

Satz 6 (Stirling). *Die Fakultät hat das asymptotische Verhalten*

$$n! \sim \sqrt{2\pi n} \left(\frac{n}{e}\right)^n.$$

Beweis. Wir bezeichnen mit $\varphi: \mathbb{R} \rightarrow \mathbb{R}$ die wie folgt definierte Funktion:

$$\varphi(x) := \frac{1}{2}x(1-x) \quad \text{für } x \in [0, 1],$$

$$\varphi(x+n) := \varphi(x) \quad \text{für alle } n \in \mathbb{Z} \text{ und } x \in [0, 1].$$

Aus der Trapez-Regel (§19, Satz 7) erhalten wir wegen $\log''(x) = -1/x^2$ die Beziehung

$$\int_k^{k+1} \log x dx = \frac{1}{2}(\log(k) + \log(k+1)) + \int_k^{k+1} \frac{\varphi(x)}{x^2} dx.$$

Summation über $k = 1, \dots, n-1$ ergibt

$$\int_1^n \log x dx = \sum_{k=1}^n \log k - \frac{1}{2} \log n + \int_1^n \frac{\varphi(x)}{x^2} dx.$$

Da $\int_1^n \log x dx = n \log n - n + 1$, folgt daraus

$$\sum_{k=1}^n \log k = \left(n + \frac{1}{2}\right) \log n - n + a_n,$$

wobei

$$a_n := 1 - \int_1^n \frac{\varphi(x)}{x^2} dx.$$

Nehmen wir von beiden Seiten die Exponentialfunktion, so erhalten wir mit $c_n := e^{a_n}$

$$n! = n^{n+\frac{1}{2}} e^{-n} c_n, \quad \text{also } c_n = \frac{n!}{\sqrt{n}} \frac{e^n}{n^n}.$$

Da φ beschränkt ist und $\int_1^\infty x^{-2} dx < \infty$, existiert der Grenzwert

$$a := \lim_{n \rightarrow \infty} a_n = 1 - \int_1^\infty \frac{\varphi(x)}{x^2} dx,$$

also auch der Grenzwert $c := \lim_{n \rightarrow \infty} c_n = e^a$. Es ist

$$\frac{c_n^2}{c_{2n}} = \frac{(n!)^2 \sqrt{2n} (2n)^{2n}}{n^{2n+1} (2n)!} = \sqrt{2} \frac{2^{2n} (n!)^2}{\sqrt{n} (2n)!}$$

und $\lim_{n \rightarrow \infty} \frac{c_n^2}{c_{2n}} = \frac{c^2}{c} = c$. Um c zu berechnen, benützen wir das Wallis'sche Produkt (19.24)

$$\pi = 2 \prod_{k=1}^{\infty} \frac{4k^2}{4k^2 - 1} = 2 \lim_{n \rightarrow \infty} \frac{2 \cdot 2 \cdot 4 \cdot 4 \cdot \dots \cdot 2n \cdot 2n}{1 \cdot 3 \cdot 3 \cdot 5 \cdot \dots \cdot (2n-1)(2n+1)}.$$

Es gilt

$$\begin{aligned} \left(2 \prod_{k=1}^n \frac{4k^2}{4k^2 - 1} \right)^{1/2} &= \sqrt{2} \frac{2 \cdot 4 \cdot \dots \cdot 2n}{3 \cdot 5 \cdot \dots \cdot (2n-1) \sqrt{2n+1}} \\ &= \frac{1}{\sqrt{n + \frac{1}{2}}} \cdot \frac{2^2 \cdot 4^2 \cdot \dots \cdot (2n)^2}{2 \cdot 3 \cdot 4 \cdot 5 \cdot \dots \cdot (2n-1) \cdot 2n} \\ &= \frac{1}{\sqrt{n + \frac{1}{2}}} \cdot \frac{2^{2n} (n!)^2}{(2n)!}, \end{aligned}$$

also

$$\sqrt{\pi} = \lim_{n \rightarrow \infty} \frac{2^{2n} (n!)^2}{\sqrt{n} (2n)!}.$$

Daraus folgt $c = \sqrt{2\pi}$, d.h.

$$\lim_{n \rightarrow \infty} \frac{n!}{\sqrt{2\pi n} \cdot n^n e^{-n}} = 1, \quad \text{q.e.d.}$$

Zusatz zu Satz 6 (Fehlerabschätzung). *Die Fakultät von n liegt zwischen den Schranken*

$$\sqrt{2\pi n} \left(\frac{n}{e}\right)^n < n! \leq \sqrt{2\pi n} \left(\frac{n}{e}\right)^n e^{1/12n}.$$

Beweis. Mit den Bezeichnungen des Beweises von Satz 6 gilt für $n \geq 1$

$$n! = \sqrt{2\pi n} \left(\frac{n}{e}\right)^n e^{a_n - a}$$

mit $a_n - a = \int_n^\infty \frac{\varphi(x)}{x^2} dx$, wobei $\varphi(x) = \frac{1}{2}z(1-z)$ für $z = x - [x]$.

Wir müssen also zeigen

$$0 < \int_n^\infty \frac{\varphi(x)}{x^2} \leq \frac{1}{12n}.$$

Da $\varphi(x) \geq 0$, ist die erste Ungleichung klar. Um das Integral nach oben abzuschätzen, benützen wir folgenden Hilfssatz.

Hilfssatz. *Sei $f: [0, 1] \rightarrow \mathbb{R}$ eine zweimal differenzierbare konvexe Funktion. Dann gilt*

$$\int_0^1 \frac{1}{2}x(1-x)f(x) dx \leq \frac{1}{12} \int_0^1 f(x) dx.$$

Beweis. Um die Symmetrie der Funktion $\frac{1}{2}x(1-x)$ um den Punkt $\frac{1}{2}$ besser ausnützen zu können, machen wir die Substitution $t = x - \frac{1}{2}$ und setzen $g(t) := f(t + \frac{1}{2})$. Dann ist g ebenfalls konvex und es ist zu zeigen

$$\int_{-1/2}^{1/2} \left(\frac{1}{8} - \frac{1}{2}t^2\right)g(t) dt \leq \frac{1}{12} \int_{-1/2}^{1/2} g(t) dt.$$

Da $\frac{1}{8} - \frac{1}{12} = \frac{1}{24}$, ist dies gleichbedeutend mit

$$\int_{-1/2}^{1/2} \left(\frac{1}{24} - \frac{1}{2}t^2\right)g(t) dt \leq 0.$$

Dies zeigen wir mit partieller Integration. Für die Funktion

$$\Psi(t) := \frac{1}{24}t - \frac{1}{6}t^3$$

gilt $\psi'(t) = \frac{1}{24} - \frac{1}{2}t^2$ und $\psi(-\frac{1}{2}) = \psi(\frac{1}{2}) = 0$, also

$$\int_{-1/2}^{1/2} (\frac{1}{24} - \frac{1}{2}t^2)g(t) dt = - \int_{-1/2}^{1/2} \psi(t)g'(t) dt.$$

Ausnutzung der Antisymmetrie $\psi(-t) = -\psi(t)$ liefert weiter

$$\int_{-1/2}^{1/2} \psi(t)g'(t) dt = \int_0^{1/2} \psi(t)(g'(t) - g'(-t)) dt.$$

Da g konvex ist, ist g' monoton steigend, also $g'(t) - g'(-t) \geq 0$ für $t \in [0, \frac{1}{2}]$. Da außerdem $\psi(t) \geq 0$ für $t \in [0, \frac{1}{2}]$, ist das letzte Integral nicht-negativ. Daraus folgt die Behauptung des Hilfssatzes.

Nun können wir den Beweis der Fehlerabschätzung zu Ende führen. Da die Funktion $x \mapsto 1/x^2$ konvex ist, erhalten wir mit dem Hilfssatz

$$\int_n^\infty \frac{\varphi(x)}{x^2} dx \leq \frac{1}{12} \int_n^\infty \frac{dx}{x^2} = \frac{1}{12n}, \quad \text{q.e.d.}$$

Die Fehlerabschätzung für $n!$ sagt, dass der Näherungswert $\sqrt{2\pi n} \left(\frac{n}{e}\right)^n$ zwar zu klein ist, aber der relative Fehler ist höchstens gleich $e^{1/12n} - 1$. Etwa für $n = 10$ ist der Fehler weniger als ein Prozent, für $n = 100$ weniger als ein Promille. Beispielsweise gilt mit einer Genauigkeit von 1 Promille

$$100! \approx 0.9325 \cdot 10^{158}.$$

Der exakte Wert von $100!$ kann z.B. mit dem ARIBAS-Befehl

```
==> factorial(100).
-: 933_26215_44394_41526_81699_23885_62667_00490_71596_
82643_81621_46859_29638_95217_59999_32299_15608_94146_
39761_56518_28625_36979_20827_22375_82511_85210_91686_
40000_00000_00000_00000_00000
```

ermittelt werden (wobei fraglich ist, ob für praktische Zwecke jemals der exakte Wert von $100!$ nötig ist). Die Stirlingsche Formel findet Anwendung u.a. in der Wahrscheinlichkeitstheorie und Statistik.

AUFGABEN

20.1. Man untersuche das Konvergenzverhalten der Reihen

$$\sum_{n=2}^{\infty} \frac{1}{n(\log n)^{\alpha}} \quad (\alpha \geq 0).$$

20.2. Für eine ganze Zahl $n \geq 1$ sei

$$\lambda(n) := \frac{1}{n} - \log\left(1 + \frac{1}{n}\right).$$

Man zeige:

a) $0 < \lambda(n) \leq \frac{1}{2n^2}.$

b) Für die Euler-Mascheronische Konstante γ gilt

$$\gamma = \sum_{n=1}^{\infty} \lambda(n).$$

20.3. Man beweise für $n \rightarrow \infty$

$$\begin{aligned} \sum_{k=1}^n \frac{1}{2k-1} &= 1 + \frac{1}{3} + \frac{1}{5} + \dots + \frac{1}{2n-1} \\ &= \frac{1}{2} \log n + \left(\frac{1}{2}\gamma + \log 2\right) + o(1). \end{aligned}$$

Dabei ist γ die Euler-Mascheronische Konstante.

20.4. Man zeige, dass folgende Umordnung der alternierenden harmonischen Reihe gegen $\frac{3}{2} \log 2$ konvergiert.

$$\begin{aligned} 1 + \frac{1}{3} - \frac{1}{2} + \frac{1}{5} + \frac{1}{7} - \frac{1}{4} + \frac{1}{9} + \frac{1}{11} - \frac{1}{6} + \frac{1}{13} + \frac{1}{15} - \frac{1}{8} + \dots \\ = \sum_{k=1}^{\infty} \left(\frac{1}{4k-3} + \frac{1}{4k-1} - \frac{1}{2k} \right) = \frac{3}{2} \log 2. \end{aligned}$$

20.5. Man zeige: Es gibt eine Konstante $\beta \in [0, 1]$, so dass

$$\sum_{k=2}^n \frac{1}{k \log k} = \log \log n + \beta + o(1) \quad \text{für } n \rightarrow \infty.$$

20.6. Für welche $\alpha \in \mathbb{R}$ konvergieren die folgenden uneigentlichen Integrale:

i)
$$\int_0^{\infty} \frac{\sin x}{x^\alpha} dx,$$

ii)
$$\int_0^{\infty} \sin(x^\alpha) dx,$$

iii)
$$\int_0^{\infty} \frac{\cos x}{x^\alpha} dx,$$

iv)
$$\int_0^{\infty} \cos(x^\alpha) dx.$$

20.7. Man beweise die asymptotische Beziehung

$$\frac{1}{2^{2n}} \binom{2n}{n} \sim \frac{1}{\sqrt{\pi n}}.$$

Bemerkung. Die Zahl $\frac{1}{2^{2n}} \binom{2n}{n}$ kann interpretiert werden als die Wahrscheinlichkeit dafür, dass beim $2n$ -maligen unabhängigen Werfen einer Münze genau n -mal das Ergebnis ‘Zahl’ auftritt.

20.8. Der Definitionsbereich der Gamma-Funktion kann wie folgt von $\mathbb{R}_+^*$ auf $D := \{t \in \mathbb{R} : -t \notin \mathbb{N}\}$ erweitert werden: Für negatives nicht-ganzes x wähle man eine natürliche Zahl n , so dass $x + n + 1 > 0$ und setze

$$\Gamma(x) := \frac{\Gamma(x + n + 1)}{x(x + 1) \cdots (x + n)}.$$

Man zeige, dass diese Definition unabhängig von der Wahl von n ist und damit die Produktdarstellung der Gamma-Funktion (Corollar zu Satz 5) für alle $x \in D$ gilt.

20.9. Man beweise für $x > 0$ die Formel

$$\Gamma\left(\frac{x}{2}\right) \Gamma\left(\frac{x+1}{2}\right) = 2^{1-x} \sqrt{\pi} \Gamma(x).$$

Anleitung. Man zeige, dass die Funktion $F(x) := 2^x \Gamma(\frac{x}{2}) \Gamma(\frac{x+1}{2})$ der Funktionalgleichung $xF(x) = F(x+1)$ genügt und logarithmisch konvex ist.

20.10. Die Eulersche Beta-Funktion ist für $x, y \in \mathbb{R}_+^*$ definiert durch

$$B(x, y) := \int_0^1 t^{x-1} (1-t)^{y-1} dt.$$

a) Man zeige, dass dieses uneigentliche Integral konvergiert.

b) Man beweise: Für festes $y > 0$ ist die Funktion $x \mapsto B(x, y)$ auf $\mathbb{R}_+^*$ logarithmisch konvex und genügt der Funktionalgleichung

$$xB(x, y) = (x+y)B(x+1, y).$$

c) Man beweise die Formel

$$B(x, y) = \frac{\Gamma(x)\Gamma(y)}{\Gamma(x+y)} \quad \text{für alle } x, y > 0.$$

Anleitung. Betrachte (für festes y) die Funktion $x \mapsto B(x, y)\Gamma(x+y)/\Gamma(y)$.

20.11. Man beweise:

a) Für alle $\alpha \in \mathbb{R}$ mit $0 < \alpha < 1$ gilt

$$\int_0^\infty \frac{x^{\alpha-1}}{1+x} dx = B(\alpha, 1-\alpha).$$

Bemerkung. In §21, Corollar zu Satz 7, wird bewiesen, dass

$$B(\alpha, 1-\alpha) = \frac{\pi}{\sin \pi \alpha}.$$

b) Für jede natürliche Zahl $n \geq 2$ gilt

$$\int_0^\infty \frac{dx}{1+x^n} = \frac{1}{n} B\left(\frac{1}{n}, 1 - \frac{1}{n}\right).$$

§ 21 Gleichmäßige Konvergenz von Funktionenfolgen

Der Begriff der Konvergenz einer Folge von Funktionen (f_n) gegen eine Funktion f , die alle denselben Definitionsbereich D haben, kann einfach auf den Konvergenzbegriff für Zahlenfolgen zurückgeführt werden: Man verlangt, dass an jeder Stelle $x \in D$ die Zahlenfolge $f_n(x)$, für $n \rightarrow \infty$ gegen $f(x)$ konvergiert. Wenn man Aussagen über die Funktion f aufgrund der Eigenschaften der Funktionen f_n beweisen will, reicht jedoch meistens diese so genannte punktweise Konvergenz nicht aus. Man braucht zusätzlich, dass die Konvergenz gleichmäßig ist, das heißt grob gesprochen, dass die Konvergenz der Folge $(f_n(x))$ gegen $f(x)$ für alle $x \in D$ gleich schnell ist. Beispielsweise gilt bei gleichmäßiger Konvergenz, dass die Grenzfunktion f wieder stetig ist, falls alle f_n stetig sind. Die gleichmäßige Konvergenz spielt auch bei der Frage eine Rolle, wann Differentiation und Integration von Funktionen mit der Limesbildung vertauschbar sind. Besonders wichtige Beispiele für gleichmäßig konvergente Funktionenfolgen liefern die Partialsummen von Potenzreihen.

Definition. Sei K eine Menge und seien $f_n: K \rightarrow \mathbb{C}$, $n \in \mathbb{N}$, Funktionen.

- a) Die Folge (f_n) konvergiert *punktweise* gegen eine Funktion $f: K \rightarrow \mathbb{C}$, falls für jedes $x \in K$ die Folge $(f_n(x))$ gegen $f(x)$ konvergiert, d.h. wenn gilt:

Zu jedem $x \in K$ und $\varepsilon > 0$ existiert ein $N = N(x, \varepsilon)$, so dass
 $|f_n(x) - f(x)| < \varepsilon$ für alle $n \geq N$.

- b) Die Folge (f_n) konvergiert *gleichmäßig* gegen eine Funktion $f: K \rightarrow \mathbb{C}$, falls gilt:

Zu jedem $\varepsilon > 0$ existiert ein $N = N(\varepsilon)$, so dass
 $|f_n(x) - f(x)| < \varepsilon$ für alle $x \in K$ und alle $n \geq N$.

Der Unterschied ist also der, dass im Fall gleichmäßiger Konvergenz N nur von ε , nicht aber von x abhängt. Konvergiert eine Funktionenfolge gleichmäßig, so auch punktweise. Die Umkehrung gilt jedoch nicht, wie folgendes Beispiel zeigt:

(21.1) Für $n \geq 2$ sei $f_n: [0, 1] \rightarrow \mathbb{R}$ definiert durch

$$f_n(x) := \max\left(n - n^2\left|x - \frac{1}{n}\right|, 0\right) \quad (\text{Bild 21.1}).$$

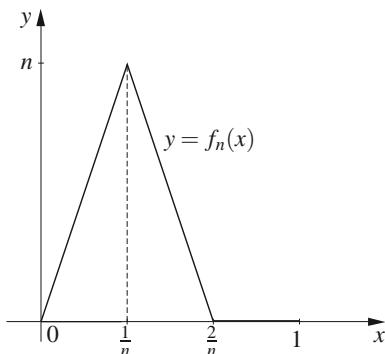


Bild 21.1

Wir zeigen, dass die Folge (f_n) punktweise gegen 0 konvergiert.

1. Für $x = 0$ ist $f_n(x) = 0$ für alle n .
2. Zu jedem $x \in]0, 1]$ existiert ein $N \geq 2$, so dass

$$\frac{2}{n} \leq x \quad \text{für alle } n \geq N.$$

Damit gilt $f_n(x) = 0$ für alle $n \geq N$, d.h. $\lim_{n \rightarrow \infty} f_n(x) = 0$.

Die Folge (f_n) konvergiert jedoch nicht gleichmäßig gegen 0, denn für kein $n \geq 2$ gilt

$$|f_n(x) - 0| < 1 \quad \text{für alle } x \in [0, 1].$$

Stetigkeit und gleichmäßige Konvergenz

Satz 1. Sei $K \subset \mathbb{C}$ und $f_n: K \rightarrow \mathbb{C}$, $n \in \mathbb{N}$, eine Folge stetiger Funktionen, die gleichmäßig gegen die Funktion $f: K \rightarrow \mathbb{C}$ konvergiere. Dann ist auch f stetig.

Anders ausgedrückt: Der Limes einer gleichmäßig konvergenten Folge stetiger Funktionen ist wieder stetig.

Beweis. Sei $x \in K$. Es ist zu zeigen, dass es zu jedem $\varepsilon > 0$ ein $\delta > 0$ gibt, so dass

$$|f(x) - f(x')| < \varepsilon \quad \text{für alle } x' \in K \text{ mit } |x - x'| < \delta.$$

Da die Folge (f_n) gleichmäßig gegen f konvergiert, existiert ein $N \in \mathbb{N}$, so dass

$$|f_N(\xi) - f(\xi)| < \frac{\varepsilon}{3} \quad \text{für alle } \xi \in K.$$

Da f_N im Punkt x stetig ist, existiert ein $\delta > 0$, so dass

$$|f_N(x) - f_N(x')| < \frac{\varepsilon}{3} \quad \text{für alle } x' \in K \text{ mit } |x - x'| < \delta.$$

Daher gilt für alle $x' \in K$ mit $|x - x'| < \delta$

$$\begin{aligned} |f(x) - f(x')| &\leq |f(x) - f_N(x)| + |f_N(x) - f_N(x')| + |f_N(x') - f(x')| \\ &< \frac{\varepsilon}{3} + \frac{\varepsilon}{3} + \frac{\varepsilon}{3} = \varepsilon, \quad \text{q.e.d.} \end{aligned}$$

Bemerkung. Konvergiert eine Folge stetiger Funktionen nur punktweise, so braucht die Grenzfunktion nicht stetig zu sein. Dazu betrachten wir folgendes Beispiel.

(21.2) Sägezahnfunktion. Sei $\sigma: \mathbb{R} \rightarrow \mathbb{R}$ wie folgt definiert (Bild 21.2):

$$\sigma(0) := 0$$

$$\sigma(x) := \frac{\pi - x}{2} \quad \text{für } x \in]0, 2\pi[,$$

$$\sigma(x + 2n\pi) := \sigma(x) \quad \text{für } n \in \mathbb{Z} \text{ und } x \in [0, 2\pi[.$$

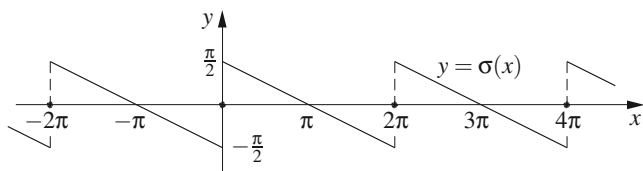


Bild 21.2

Nach Beispiel (19.25) gilt

$$\sigma(x) = \sum_{k=1}^{\infty} \frac{\sin kx}{k} \quad \text{für alle } x \in \mathbb{R}.$$

Wir hatten in (19.25) diese Beziehung für $0 < x < 2\pi$ bewiesen; für $x = 0$ gilt sie trivialerweise und für $2n\pi \leq x < 2(n+1)\pi$ folgt sie daraus, dass $\sin k(x + 2n\pi) = \sin kx$.

Die Partialsummen der Reihe sind stetig auf ganz $\mathbb{R}$, der Limes jedoch unstetig an den Stellen $x = 2n\pi$, ($n \in \mathbb{Z}$). Also kann die Reihe auf $\mathbb{R}$ nicht gleichmäßig konvergieren. Wir wollen jedoch zeigen, dass die Reihe für jedes $\delta \in]0, \pi[$ im Intervall $[\delta, 2\pi - \delta]$ gleichmäßig konvergiert. Dazu setzen wir

$$s_n(x) := \sum_{k=1}^n \sin kx = \operatorname{Im} \left(\sum_{k=1}^n e^{ikx} \right).$$

Für $\delta \leq x \leq 2\pi - \delta$ gilt

$$|s_n(x)| \leq \left| \sum_{k=1}^n e^{ikx} \right| = \left| \frac{e^{inx} - 1}{e^{ix} - 1} \right| \leq \frac{2}{|e^{ix/2} - e^{-ix/2}|} = \frac{1}{\sin \frac{x}{2}} \leq \frac{1}{\sin \frac{\delta}{2}}.$$

Es folgt für $m \geq n > 0$

$$\begin{aligned} \left| \sum_{k=n}^m \frac{\sin kx}{k} \right| &= \left| \sum_{k=n}^m \frac{s_k(x) - s_{k-1}(x)}{k} \right| \\ &= \left| \sum_{k=n}^m s_k(x) \left(\frac{1}{k} - \frac{1}{k+1} \right) + \frac{s_m(x)}{m+1} - \frac{s_{n-1}(x)}{n} \right| \\ &\leq \frac{1}{\sin \frac{\delta}{2}} \left(\frac{1}{n} - \frac{1}{m+1} + \frac{1}{m+1} + \frac{1}{n} \right) = \frac{2}{n \sin \frac{\delta}{2}}, \end{aligned}$$

also auch

$$\left| \sum_{k=n}^{\infty} \frac{\sin kx}{k} \right| \leq \frac{2}{n \sin \frac{\delta}{2}} \quad \text{für alle } x \in [\delta, 2\pi - \delta].$$

Daraus folgt die behauptete gleichmäßige Konvergenz. Gemäß Satz 1 ist die Summe der Reihe im Intervall $[\delta, 2\pi - \delta]$ stetig. Aber natürlich kann der Limes einer Folge stetiger Funktionen auch stetig sein, wenn die Konvergenz nicht gleichmäßig, sondern nur punktweise ist, siehe Beispiel (21.1).

Definition (Supremumsnorm). Sei K eine Menge und $f: K \rightarrow \mathbb{C}$ eine Funktion. Dann setzt man

$$\|f\|_K := \sup\{|f(x)| : x \in K\}.$$

Bemerkung. Es gilt $\|f\|_K \in \mathbb{R}_+ \cup \{\infty\}$. Die Funktion f ist genau dann beschränkt, wenn $\|f\|_K < \infty$, d.h. $\|f\|_K \in \mathbb{R}_+$. Sind Missverständnisse ausgeschlossen, schreibt man oft kurz $\|f\|$ statt $\|f\|_K$.

Mithilfe der Supremumsnorm läßt sich die Definition der gleichmäßigen Konvergenz so umformen: Eine Folge $f_n: K \rightarrow \mathbb{C}$, $n \in \mathbb{N}$, von Funktionen konvergiert genau dann gleichmäßig auf K gegen $f: K \rightarrow \mathbb{C}$, wenn

$$\lim_{n \rightarrow \infty} \|f_n - f\|_K = 0.$$

Denn die Bedingung $\|f_n - f\|_K \leq \varepsilon$ ist gleichbedeutend mit $|f_n(x) - f(x)| \leq \varepsilon$ für alle $x \in K$. Die Bedingung $\|f_n - f\|_K \leq \varepsilon$ bedeutet im Fall reeller Funktionen, dass der Graph von f_n ganz im “ ε -Streifen” zwischen $f - \varepsilon$ und $f + \varepsilon$ liegt (Bild 21.3).

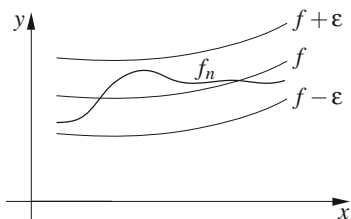


Bild 21.3

Satz 2 (Konvergenzkriterium von Weierstraß). Seien $f_n: K \rightarrow \mathbb{C}$, $n \in \mathbb{N}$, Funktionen. Es gelte

$$\sum_{n=0}^{\infty} \|f_n\|_K < \infty.$$

Dann konvergiert die Reihe $\sum_{n=0}^{\infty} f_n$ absolut und gleichmäßig auf K gegen eine Funktion $F: K \rightarrow \mathbb{C}$.

Beweis. a) Wir zeigen zunächst, dass $\sum f_n$ punktweise gegen eine gewisse Funktion $F: K \rightarrow \mathbb{C}$ konvergiert.

Sei $x \in K$. Da $|f_n(x)| \leq \|f_n\|_K$, konvergiert (nach dem Majoranten-Kriterium) die Reihe $\sum f_n(x)$ absolut. Wir setzen

$$F(x) := \sum_{n=0}^{\infty} f_n(x).$$

Damit ist eine Funktion $F: K \rightarrow \mathbb{C}$ definiert.

b) Sei $F_n := \sum_{k=0}^n f_k$. Wir beweisen jetzt, dass die Folge (F_n) gleichmäßig gegen F konvergiert.

Sei $\varepsilon > 0$ vorgegeben. Aus der Konvergenz von $\sum \|f_n\|_K$ folgt, dass es ein N gibt, so dass

$$\sum_{k=n+1}^{\infty} \|f_k\|_K < \varepsilon \quad \text{für alle } n \geq N.$$

Dann gilt für $n \geq N$ und alle $x \in K$

$$|F_n(x) - F(x)| = \left| \sum_{k=n+1}^{\infty} f_k(x) \right| \leq \sum_{k=n+1}^{\infty} |f_k(x)| \leq \sum_{k=n+1}^{\infty} \|f_k\|_K < \varepsilon.$$

(21.3) Die Reihe $\sum_{n=1}^{\infty} \frac{\cos nx}{n^2}$ konvergiert gleichmäßig auf $\mathbb{R}$, denn für

$$f_n(x) := \frac{\cos nx}{n^2} \quad \text{gilt} \quad \|f_n\|_{\mathbb{R}} = \frac{1}{n^2} \quad \text{und} \quad \sum_{n=1}^{\infty} \frac{1}{n^2} < \infty.$$

Potenzreihen

Besonders gute Konvergenz-Eigenschaften haben die Potenzreihen.

Satz 3. Sei $(c_n)_{n \in \mathbb{N}}$ eine Folge komplexer Zahlen und $a \in \mathbb{C}$. Die Potenzreihe

$$f(z) = \sum_{n=0}^{\infty} c_n (z-a)^n$$

konvergiere für ein $z_1 \in \mathbb{C}$, $z_1 \neq a$. Sei r eine reelle Zahl mit $0 < r < |z_1 - a|$ und

$$K(a, r) := \{z \in \mathbb{C} : |z - a| \leq r\} \quad (\text{Bild 21.4}).$$

Dann konvergiert die Potenzreihe absolut und gleichmäßig auf $K(a, r)$. Die formal differenzierte Potenzreihe

$$g(z) = \sum_{n=1}^{\infty} n c_n (z-a)^{n-1}$$

konvergiert ebenfalls absolut und gleichmäßig auf $K(a, r)$.

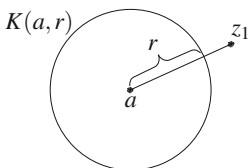


Bild 21.4

Beweis

a) Sei $f_n(z) := c_n(z-a)^n$, also $f = \sum_{n=0}^{\infty} f_n$. Da $\sum_{n=0}^{\infty} f_n(z_1)$ nach Voraussetzung konvergiert, existiert ein $M \in \mathbb{R}_+$, so dass $|f_n(z_1)| \leq M$ für alle $n \in \mathbb{N}$. Für alle $z \in K(a, r)$ gilt dann

$$|f_n(z)| = |c_n(z-a)^n| = |c_n(z_1-a)^n| \cdot \left| \frac{z-a}{z_1-a} \right|^n \leq M \theta^n,$$

wobei

$$\theta := \frac{r}{|z_1 - a|} \in]0, 1[.$$

Es gilt also $\|f_n\|_{K(a,r)} \leq M\theta^n$. Da $\sum_{n=0}^{\infty} M\theta^n$ konvergiert (geometrische Reihe), konvergiert $\sum_{n=0}^{\infty} f_n$ nach Satz 2 absolut und gleichmäßig auf $K(a, r)$.

b) Sei $g_n(z) := nc_n(z-a)^{n-1}$, also $g = \sum g_n$. Wie unter a) zeigt man, dass

$$\|g_n\|_{K(a,r)} \leq nM\theta^{n-1}.$$

Nach dem Quotienten-Kriterium konvergiert $\sum_{n=1}^{\infty} nM\theta^{n-1}$, also folgt aus Satz 2 die Behauptung.

Definition. Sei $f(z) = \sum_{n=0}^{\infty} c_n(z-a)^n$ eine Potenzreihe. Dann heißt

$$R := \sup \left\{ |z-a| : \sum_{n=0}^{\infty} c_n(z-a)^n \text{ konvergiert} \right\}$$

Konvergenzradius der Potenzreihe.

Bemerkung. Es gilt $R \in \mathbb{R}_+ \cup \{\infty\}$. Nach Satz 3 konvergiert für jedes $r \in [0, R[$ die Potenzreihe gleichmäßig auf $K(a, r)$. Die Potenzreihe konvergiert sogar in der offenen Kreisscheibe

$$B(a, R) := \{z \in \mathbb{C} : |z-a| < R\},$$

da $B(a, R) = \bigcup_{r < R} K(a, r)$. In $B(a, R)$ ist die Konvergenz i.Allg. jedoch nicht gleichmäßig. Aus Satz 1 folgt: Der Limes einer Potenzreihe stellt eine im Innern des Konvergenzkreises stetige Funktion dar.

Corollar (Identitätssatz für Potenzreihen). *Seien*

$$f(z) = \sum_{n=0}^{\infty} c_n(z-a)^n \quad \text{und} \quad g(z) = \sum_{n=0}^{\infty} \gamma_n(z-a)^n$$

zwei Potenzreihen mit komplexen Koeffizienten, die beide (mindestens) im Kreis $B(a, r)$, $r > 0$, konvergieren. Es gebe eine Folge von Punkten $z_v \in B(a, r) \setminus \{a\}$, $v \in \mathbb{N}$, mit

$$f(z_v) = g(z_v) \quad \text{für alle } v \in \mathbb{N} \quad \text{und} \quad \lim_{v \rightarrow \infty} z_v = a.$$

Dann sind f und g identisch, d.h. $c_n = \gamma_n$ für alle $n \in \mathbb{N}$.

Beweis. O.B.d.A. können wir annehmen, dass g identisch null ist (andernfalls betrachte man die Differenz $f - g$). Es gilt dann $f(z_v) = 0$ für alle v und wir müssen zeigen, dass alle Koeffizienten c_n verschwinden. Angenommen, dies

sei nicht der Fall. Sei $m \geq 0$ der kleinste Index mit $c_m \neq 0$. Wir betrachten die Potenzreihe

$$F(z) := \frac{f(z)}{(z-a)^m} = \sum_{k=0}^{\infty} c_{m+k}(z-a)^k.$$

F konvergiert ebenfalls in $B(a, r)$ und stellt dort eine stetige Funktion dar. Nach Voraussetzung gilt $F(z_v) = 0$ für alle v und aus der Stetigkeit folgt

$$0 = \lim_{v \rightarrow \infty} F(z_v) = F(a) = c_m, \quad \text{Widerspruch!}$$

Also sind doch alle Koeffizienten $c_n = 0$ und der Identitätssatz ist bewiesen.

Bemerkung. Ist $f(x) = \sum c_n(x-a)^n$ eine reelle Potenzreihe, die für ein reelles $x_1 \neq a$ konvergiert, so folgt aus Satz 3, dass die Potenzreihe automatisch auch in einem Kreis in der komplexen Ebene konvergiert. Reelle Funktionen, die durch Potenzreihen dargestellt werden, können so “ins Komplexe” fortgesetzt werden. Die systematische Untersuchung der durch Potenzreihen darstellbaren Funktionen ist Gegenstand der so genannten Funktionentheorie [FB], [FL].

Integration und Limesbildung

Wir betrachten jetzt die folgende Situation: Gegeben sei eine Folge (f_n) auf einem Intervall I definierter Funktionen, die gegen eine Funktion f konvergiert. Die Integrale der Funktionen f_n über das Intervall seien bekannt. Was kann man dann über das Integral der Grenzfunktion f schließen? Bei gleichmäßiger Konvergenz hat man folgende Aussage:

Satz 4. Sei $f_n: [a, b] \rightarrow \mathbb{R}$, $n \in \mathbb{N}$, eine Folge stetiger Funktionen. Die Folge konvergiere auf $[a, b]$ gleichmäßig gegen die Funktion $f: [a, b] \rightarrow \mathbb{R}$. Dann gilt

$$\int_a^b f(x) dx = \lim_{n \rightarrow \infty} \int_a^b f_n(x) dx.$$

Bemerkung. Dieser Satz sagt also, dass man bei gleichmäßiger Konvergenz Integration und Limesbildung „vertauschen“ darf.

Beweis. Nach Satz 1 ist f wieder stetig, also integrierbar. Es gilt

$$\left| \int_a^b f(x) dx - \int_a^b f_n(x) dx \right| \leq \int_a^b |f(x) - f_n(x)| dx \leq (b-a) \|f - f_n\|.$$

Dies konvergiert für $n \rightarrow \infty$ gegen null, q.e.d.

Beispiele

(21.4) Satz 4 gilt nicht für punktweise Konvergenz, wie die in (21.1) definierten Funktionen $f_n: [0, 1] \rightarrow \mathbb{R}$ zeigen. Es ist nämlich

$$\int_0^1 f_n(x) dx = 1 \quad \text{für alle } n \geq 2,$$

aber

$$\int_0^1 (\lim f_n(x)) dx = \int_0^1 0 dx = 0.$$

(21.5) Nach Satz 3 konvergiert die Exponentialreihe gleichmäßig auf jedem Intervall $[a, b]$. Also gilt

$$\begin{aligned} \int_a^b \exp(x) dx &= \int_a^b \left(\sum_{n=0}^{\infty} \frac{x^n}{n!} \right) dx = \sum_{n=0}^{\infty} \int_a^b \frac{x^n}{n!} dx \\ &= \sum_{n=0}^{\infty} \frac{x^{n+1}}{(n+1)!} \Big|_a^b = \sum_{n=1}^{\infty} \frac{x^n}{n!} \Big|_a^b = \exp(x) \Big|_a^b. \end{aligned}$$

Dies Resultat ist uns natürlich schon aus (19.5) bekannt.

(21.6) Als weiteres Beispiel leiten wir die Potenzreihen-Entwicklung der Funktion Integral-Sinus ab, vgl. Beispiel (20.6):

$$\text{Si}(x) = \int_0^x \frac{\sin t}{t} dt.$$

Aus der Potenzreihen-Entwicklung der Funktion $\sin$ erhält man

$$\frac{\sin t}{t} = \sum_{k=0}^{\infty} (-1)^k \frac{t^{2k}}{(2k+1)!}.$$

Diese Reihe konvergiert auf ganz $\mathbb{R}$, also gleichmäßig auf jedem kompakten Intervall $[0, x]$, ($x \geq 0$). Deshalb darf man gliedweise integrieren und erhält

$$\text{Si}(x) = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k+1}}{(2k+1)(2k+1)!} = x - \frac{x^3}{18} + \frac{x^5}{600} - \frac{x^7}{35280} \pm \dots$$

Auch diese Reihe konvergiert auf ganz $\mathbb{R}$. Z.B. ergibt sich für das erste lokale Maximum von Si an der Stelle $x = \pi$ bei Berücksichtigung aller Terme bis x^{15}

$$\text{Si}(\pi) = 1.8519370\dots = 1.1789797\dots \cdot \frac{\pi}{2}.$$

Für große Argumente x ist die Reihe zur numerischen Rechnung weniger geeignet.

(21.7) Will man Satz 4 auf uneigentliche Integrale anwenden, sind zusätzliche Überlegungen notwendig (vgl. das Gegenbeispiel in Aufgabe 21.1). Wir beweisen hier als Beispiel die Formel

$$\int_0^{\infty} \frac{x^{s-1}}{e^x - 1} dx = \Gamma(s) \zeta(s) \quad \text{für } s > 1.$$

Beweis. Dass das uneigentliche Integral $\int_0^{\infty} \frac{x^{s-1}}{e^x - 1} dx$ für $s > 1$ konvergiert, beweist man ähnlich wie bei der Gamma-Funktion durch folgende zwei Abschätzungen:

(i) Da $\lim_{x \rightarrow 0} \frac{e^x - 1}{x} = 1$, gilt

$$\frac{x^{s-1}}{e^x - 1} \leq 2x^{s-2} \quad \text{für } 0 < x \leq x_0, \quad (x_0 > 0 \text{ geeignet}).$$

(ii) Da e^x für $x \rightarrow \infty$ schneller als jede Potenz von x gegen ∞ strebt, folgt

$$\frac{x^{s-1}}{e^x - 1} \leq \frac{1}{x^2} \quad \text{für } x \geq x_1.$$

Sei nun $0 < \delta < R < \infty$. Dann gilt im Intervall $[\delta, R]$

$$\frac{x^{s-1}}{e^x - 1} = x^{s-1} e^{-x} \frac{1}{1 - e^{-x}} = x^{s-1} e^{-x} \sum_{n=0}^{\infty} e^{-nx} = \sum_{n=1}^{\infty} x^{s-1} e^{-nx},$$

wobei wegen $|e^{-x}| \leq e^{-\delta} < 1$ gleichmäßige Konvergenz vorliegt. Wir setzen zur Abkürzung

$$F(x) := \frac{x^{s-1}}{e^x - 1} \quad \text{und} \quad f_n(x) := x^{s-1} e^{-nx}.$$

Alle diese Funktionen sind positiv für $x \in \mathbb{R}_+^*$. Aus Satz 4 folgt

$$(*) \quad \int_{\delta}^R F(x) dx = \sum_{n=1}^{\infty} \int_{\delta}^R f_n(x) dx.$$

Aus (*) folgt für jedes $N \geq 1$

$$\sum_{n=1}^N \int_{\delta}^R f_n(x) dx \leq \int_0^{\infty} F(x) dx,$$

also auch (durch Grenzübergang $\delta \rightarrow 0, R \rightarrow \infty$)

$$\sum_{n=1}^N \int_0^{\infty} f_n(x) dx \leq \int_0^{\infty} F(x) dx,$$

und weiter ($N \rightarrow \infty$)

$$\sum_{n=1}^{\infty} \int_0^{\infty} f_n(x) dx \leq \int_0^{\infty} F(x) dx.$$

Andrerseits ist nach (*)

$$\int_{\delta}^R F(x) dx \leq \sum_{n=1}^{\infty} \int_0^{\infty} f_n(x) dx,$$

also auch

$$\int_0^{\infty} F(x) dx \leq \sum_{n=1}^{\infty} \int_0^{\infty} f_n(x) dx.$$

Insgesamt hat man damit die Gleichung

$$\int_0^{\infty} F(x) dx = \sum_{n=1}^{\infty} \int_0^{\infty} f_n(x) dx.$$

Wir müssen also nur noch die Integrale $\int_0^{\infty} f_n(x) dx$ ausrechnen. Mit der Substitution $t = nx$ erhält man

$$\int_{\delta}^R f_n(x) dx = \int_{\delta}^R x^{s-1} e^{-nx} dx = \frac{1}{n^s} \int_{n\delta}^{nR} t^{s-1} e^{-t} dt$$

und nach Grenzübergang $\delta \rightarrow 0, R \rightarrow \infty$

$$\int_0^{\infty} f_n(x) dx = \frac{1}{n^s} \int_0^{\infty} t^{s-1} e^{-t} dt = \frac{1}{n^s} \Gamma(s).$$

Da $\sum_{n=1}^{\infty} \frac{1}{n^s} = \zeta(s)$, folgt damit die behauptete Gleichung

$$\int_0^{\infty} \frac{x^{s-1}}{e^x - 1} dx = \zeta(s) \Gamma(s), \quad \text{q.e.d.}$$

Insbesondere folgt aus der bewiesenen Formel

$$\int_0^{\infty} \frac{dx}{x^5 (e^{1/x} - 1)} = \int_0^{\infty} \frac{t^3}{e^t - 1} dt = \Gamma(4) \zeta(4) = \frac{\pi^4}{15},$$

da $\zeta(4) = \sum_{n=1}^{\infty} \frac{1}{n^4} = \frac{\pi^4}{90}$, wie wir in §22, Satz 11 zeigen werden. Dieses Integral ist in der theoretischen Physik von Bedeutung, vgl. (17.1).

Differentiation und Limesbildung

Wir wollen uns jetzt mit der zu Satz 4 analogen Fragestellung über die Vertauschbarkeit von Differentiation und Limesbildung beschäftigen. Es stellt sich heraus, dass hier die Situation komplizierter ist. Die gleichmäßige Konvergenz der Funktionenfolge reicht nicht aus, vielmehr braucht man die gleichmäßige Konvergenz der Folge der Ableitungen.

Satz 5. *Seien $f_n: [a, b] \rightarrow \mathbb{R}$ stetig differenzierbare Funktionen ($n \in \mathbb{N}$), die punktweise gegen die Funktion $f: [a, b] \rightarrow \mathbb{R}$ konvergieren. Die Folge der Ableitungen $f'_n: [a, b] \rightarrow \mathbb{R}$ konvergiere gleichmäßig. Dann ist f differenzierbar und es gilt*

$$f'(x) = \lim_{n \rightarrow \infty} f'_n(x) \quad \text{für alle } x \in [a, b].$$

Beweis. Sei $f^* = \lim f'_n$. Nach Satz 1 ist f^* eine auf $[a, b]$ stetige Funktion. Für alle $x \in [a, b]$ gilt

$$f_n(x) = f_n(a) + \int_a^x f'_n(t) dt.$$

Nach Satz 4 konvergiert $\int_a^x f'_n(t) dt$ für $n \rightarrow \infty$ gegen $\int_a^x f^*(t) dt$, also erhält man

$$f(x) = f(a) + \int_a^x f^*(t) dt.$$

Differentiation ergibt $f'(x) = f^*(x)$, (§19, Satz 1), q.e.d.

Beispiele

(21.8) Selbst wenn (f_n) gleichmäßig gegen eine differenzierbare Funktion f konvergiert, gilt i.Allg. nicht $\lim_{n \rightarrow \infty} f'_n = f'$, wie folgendes Beispiel zeigt:

$$f_n: \mathbb{R} \rightarrow \mathbb{R}, \quad f_n(x) := \frac{1}{n} \sin nx, \quad (n \geq 1).$$

Da $\|f_n\| = \frac{1}{n}$, konvergiert die Folge (f_n) gleichmäßig gegen 0. Die Folge der Ableitungen $f'_n(x) = \cos nx$ konvergiert jedoch nicht gegen 0.

(21.9) Als Anwendung von Satz 5 berechnen wir die Summe der Reihe

$$F(x) := \sum_{n=1}^{\infty} \frac{\cos nx}{n^2},$$

die nach (21.3) gleichmäßig konvergiert. Die Reihe der Ableitungen

$$- \sum_{n=1}^{\infty} \frac{\sin nx}{n}$$

konvergiert nach (21.2) für jedes $\delta > 0$ auf dem Intervall $[\delta, 2\pi - \delta]$ gleichmäßig gegen $\frac{x-\pi}{2}$. Deshalb gilt für alle $x \in]0, 2\pi[$

$$F'(x) = \frac{x-\pi}{2}, \quad \text{d.h.} \quad F(x) = \left(\frac{x-\pi}{2}\right)^2 + c$$

mit einer Konstanten $c \in \mathbb{R}$. Da F stetig ist, gilt diese Beziehung im ganzen Intervall $[0, 2\pi]$. Um die Konstante zu bestimmen, berechnen wir das Integral

$$\int_0^{2\pi} F(x) dx = \int_0^{2\pi} \left(\frac{x-\pi}{2}\right)^2 dx + \int_0^{2\pi} c dx = \frac{\pi^3}{6} + 2\pi c.$$

Da $\int_0^{2\pi} \cos nx dx = 0$ für alle $n \geq 1$, gilt andererseits nach Satz 4

$$\int_0^{2\pi} F(x) dx = \sum_{n=1}^{\infty} \int_0^{2\pi} \frac{\cos nx}{n^2} = 0,$$

also folgt $c = -\frac{\pi^2}{12}$. Damit ist bewiesen

$$\sum_{n=1}^{\infty} \frac{\cos nx}{n^2} = \left(\frac{x-\pi}{2}\right)^2 - \frac{\pi^2}{12} \quad \text{für } 0 \leq x \leq 2\pi.$$

Insbesondere für $x = 0$ erhält man die schon in (7.2) behauptete Formel

$$\sum_{n=1}^{\infty} \frac{1}{n^2} = \frac{\pi^2}{6}.$$

Satz 6 (Ableitung von Potenzreihen). Sei $f(x) = \sum_{n=0}^{\infty} c_n(x-a)^n$ eine Potenzreihe mit dem Konvergenzradius $r > 0$, ($c_n, a \in \mathbb{R}$). Dann gilt für alle

$$x \in]a - r, a + r[$$

$$f'(x) = \sum_{n=1}^{\infty} n c_n (x-a)^{n-1}.$$

Beweis. Dies folgt unmittelbar durch Anwendung von Satz 5 auf Satz 3.

Kurz gesagt bedeutet Satz 6: Eine Potenzreihe darf gliedweise differenziert werden.

(21.10) Für $|x| < 1$ gilt

$$\sum_{n=1}^{\infty} n x^n = x \sum_{n=1}^{\infty} n x^{n-1} = x \frac{d}{dx} \sum_{n=0}^{\infty} x^n = x \frac{d}{dx} \left(\frac{1}{1-x} \right) = \frac{x}{(1-x)^2}.$$

Ein Spezialfall davon ist die Formel

$$\sum_{n=1}^{\infty} \frac{n}{2^n} = \frac{1}{2} + \frac{2}{4} + \frac{3}{8} + \frac{4}{16} + \frac{5}{32} + \frac{6}{64} + \frac{7}{128} + \cdots = 2.$$

Corollar. Die Potenzreihe

$$f(x) = \sum_{n=0}^{\infty} c_n (x-a)^n$$

konvergiere im Intervall $I :=]a - r, a + r[$, ($r > 0$). Dann ist $f: I \rightarrow \mathbb{R}$ beliebig oft differenzierbar und es gilt

$$c_n = \frac{1}{n!} f^{(n)}(a) \quad \text{für alle } n \in \mathbb{N}.$$

Beweis. Wiederholte Anwendung von Satz 6 ergibt

$$f^{(k)}(x) = \sum_{n=k}^{\infty} n(n-1) \cdots (n-k+1) c_n (x-a)^{n-k}.$$

Insbesondere folgt daraus

$$f^{(k)}(a) = k! c_k, \quad \text{d.h.} \quad c_k = \frac{1}{k!} f^{(k)}(a).$$

(21.11) Als ein Beispiel zeigen wir, dass die wie folgt definierte Funktion $f:]-\pi, \pi[\rightarrow \mathbb{R}$ unendlich oft differenzierbar ist:

$$f(x) := \frac{1}{\sin^2 x} - \frac{1}{x^2} \quad \text{für } 0 < |x| < \pi, \quad f(0) := \frac{1}{3}.$$

Beweis. Sei φ die Funktion

$$\varphi(x) := \frac{\sin x}{x} \quad \text{für } x \neq 0, \quad \varphi(0) := 1.$$

Aus der Potenzreihen-Entwicklung des Sinus folgt

$$\varphi(x) = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k}}{(2k+1)!} = 1 - \frac{x^2}{6} + \frac{x^4}{120} \pm \dots$$

Die Funktion φ ist also auf ganz $\mathbb{R}$ beliebig oft differenzierbar und im Intervall $]-\pi, \pi[$ ungleich 0. Nun lässt sich die Funktion f für $0 < |x| < \pi$ schreiben als

$$(*) \quad f(x) = \frac{1 - \varphi(x)^2}{x^2 \varphi(x)^2} = \frac{1 - \varphi(x)}{x^2} \cdot \frac{1 + \varphi(x)}{\varphi(x)^2} = \psi(x) \cdot \frac{1 + \varphi(x)}{\varphi(x)^2}$$

mit

$$\psi(x) = \frac{1 - \varphi(x)}{x^2} = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k}}{(2k+3)!} = \frac{1}{6} - \frac{x^2}{120} \pm \dots$$

Auch die Funktion ψ ist auf ganz $\mathbb{R}$ beliebig oft differenzierbar. Wegen $\psi(0) = \frac{1}{6}$ und $\varphi(0) = 1$ gilt die Darstellung $(*)$ auch für $x = 0$. Daraus folgt die Behauptung.

(21.12) Wir untersuchen jetzt die auf $\mathbb{R} \setminus \mathbb{Z}$ definierte Funktion

$$F(x) := \sum_{n \in \mathbb{Z}} \frac{1}{(x-n)^2}.$$

(Eine Summe $\sum_{n \in \mathbb{Z}} c_n$ ist als $\sum_{n=0}^{\infty} c_n + \sum_{n=1}^{\infty} c_{-n}$ zu verstehen.)

Zur Konvergenz. Sei $R > 0$ beliebig vorgegeben. Dann gilt für alle $x \in [-R, R]$ und alle $|n| \geq 2R$

$$|x - n| \geq \frac{|n|}{2}, \quad \text{also} \quad \frac{1}{(x-n)^2} \leq \frac{4}{n^2}.$$

Da $\sum_{n=1}^{\infty} \frac{1}{n^2} < \infty$, folgt mit dem Weierstraß'schen Konvergenzkriterium (Satz 2),

dass die Reihe $\sum_{|n| \geq 2R} \frac{1}{(x-n)^2}$ auf dem Intervall $[-R, R]$ absolut und gleichmäßig

konvergiert. Es folgt, dass die Reihe $\sum_{n \in \mathbb{Z}} \frac{1}{(x-n)^2}$ auf jedem kompakten Intervall $[a, b] \subset \mathbb{R}$, in dem keine ganze Zahl liegt, absolut und gleichmäßig konvergiert, also F eine in $\mathbb{R} \setminus \mathbb{Z}$ stetige Funktion darstellt.

Behauptung. Für alle $x \in \mathbb{R} \setminus \mathbb{Z}$ gilt

$$(S) \quad \sum_{n \in \mathbb{Z}} \frac{1}{(x-n)^2} = \left(\frac{\pi}{\sin \pi x} \right)^2.$$

Beweis. a) Wir zeigen zunächst, dass die linke Seite die Periode 1 hat, d.h. $F(x) = F(x+1)$ für alle $x \in \mathbb{R} \setminus \mathbb{Z}$:

$$F(x+1) = \sum_{n \in \mathbb{Z}} \frac{1}{(x+1-n)^2} = \sum_{n \in \mathbb{Z}} \frac{1}{(x-(n-1))^2} = \sum_{n \in \mathbb{Z}} \frac{1}{(x-n)^2} = F(x),$$

denn durchläuft n alle ganzen Zahlen, so auch $n-1$.

Die rechte Seite hat natürlich auch die Periode 1, denn $\sin \pi(x+1) = -\sin \pi x$.

b) Wir zeigen jetzt, dass sich die Differenz aus linker und rechter Seite stetig in alle Punkte $n \in \mathbb{Z}$ fortsetzen lässt. Wegen der Periodizität brauchen wir das nur an der Stelle 0 zu zeigen. Nach Beispiel (21.11) lässt sich $\frac{1}{\sin^2 x} - \frac{1}{x^2}$ stetig nach 0 fortsetzen (mit dem Wert $1/3$). Ersetzt man hierin die Variable x durch πx und multipliziert mit π^2 , erhält man, dass sich

$$\left(\frac{\pi}{\sin \pi x} \right)^2 - \frac{1}{x^2}$$

stetig in den Nullpunkt fortsetzen lässt (mit dem Wert $\pi^2/3$). Daraus folgt aber die Behauptung, denn

$$F(x) = \frac{1}{x^2} + \sum_{|n| \geq 1} \frac{1}{(x-n)^2},$$

und die letzte Summe ist stetig im Nullpunkt.

c) Wir leiten jetzt für die linke Seite von (S) folgende Funktionalgleichung her: Für alle $x \in \mathbb{R}$, so dass $2x$ keine ganze Zahl ist, gilt

$$4F(2x) = F(x) + F(x + \tfrac{1}{2}).$$

Dies sieht man so:

$$\begin{aligned} 4F(2x) &= \sum_{n \in \mathbb{Z}} \frac{4}{(2x-n)^2} = \sum_{n \in \mathbb{Z}} \frac{1}{(x-\frac{n}{2})^2} \\ &= \sum_{n \in \mathbb{Z}} \frac{1}{(x-n)^2} + \sum_{n \in \mathbb{Z}} \frac{1}{(x+\frac{1}{2}-n)^2} = F(x) + F(x + \tfrac{1}{2}). \end{aligned}$$

d) Die rechte Seite $G(x) := (\pi/\sin \pi x)^2$ erfüllt die zu c) analoge Funktionalgleichung. Wir benützen dazu die Formel $\sin 2\alpha = 2 \sin \alpha \cos \alpha$.

$$\begin{aligned} G(x) + G(x + \tfrac{1}{2}) &= \frac{\pi^2}{\sin^2 \pi x} + \frac{\pi^2}{\cos^2 \pi x} \\ &= \frac{\pi^2(\cos^2 \pi x + \sin^2 \pi x)}{\sin^2 \pi x \cos^2 \pi x} = \frac{4\pi^2}{\sin^2 2\pi x} = 4G(2x). \end{aligned}$$

e) Nach b) kann die Funktion $H(x) := F(x) - G(x)$ stetig auf ganz $\mathbb{R}$ fortgesetzt werden und genügt nach c) und d) der Funktionalgleichung

$$4H(2x) = H(x) + H(x + \tfrac{1}{2}) \quad \text{für alle } x \in \mathbb{R}.$$

Sei $M := \sup\{|H(x)| : 0 \leq x \leq 1\}$. Dieses Supremum wird in einem gewissen Punkt $a \in [0, 1]$ angenommen. Mit $b := a/2$ folgt aus der Funktionalgleichung

$$4M = |4H(2b)| \leq |H(b)| + |H(b + \tfrac{1}{2})| \leq 2M.$$

Diese Ungleichung kann aber nur bestehen, wenn $M = 0$, also $H(x) = 0$ für alle $x \in [0, 1]$. Wegen der Periodizität ist H auf ganz $\mathbb{R}$ identisch 0, d.h. $F \equiv G$. Damit haben wir die Formel

$$\sum_{n \in \mathbb{Z}} \frac{1}{(x-n)^2} = \left(\frac{\pi}{\sin \pi x} \right)^2 \quad \text{für alle } x \in \mathbb{R} \setminus \mathbb{Z}$$

bewiesen. Aus ihr können wir nun weitere interessante Formeln von Euler ableiten:

Satz 7 (Euler).

a) (Partialbruch-Zerlegung des Cotangens) Für alle $x \in \mathbb{R} \setminus \mathbb{Z}$ gilt

$$\pi \cot \pi x = \frac{1}{x} + \sum_{n=1}^{\infty} \frac{2x}{x^2 - n^2} = \frac{1}{x} + \sum_{n=1}^{\infty} \left(\frac{1}{x-n} + \frac{1}{x+n} \right).$$

Die Reihe konvergiert gleichmäßig auf jedem kompakten Intervall, das keinen Punkt aus $\mathbb{Z}$ enthält.

b) (Sinus-Produkt) Für alle $x \in \mathbb{R}$ gilt

$$\sin \pi x = \pi x \prod_{n=1}^{\infty} \left(1 - \frac{x^2}{n^2} \right).$$

Das unendliche Produkt konvergiert gleichmäßig auf jedem kompakten Intervall von $\mathbb{R}$.

Beweis. a) Sei $R > 0$ beliebig vorgegeben. Für $|x| \leq R$ und $n \geq 2R$ gilt dann

$$\left| \frac{2x}{x^2 - n^2} \right| < \frac{4R}{n^2},$$

woraus sich die gleichmäßige Konvergenz von $\sum_{n \geq 2R} \frac{2x}{x^2 - n^2}$ auf dem Intervall $[-R, R]$ ergibt. Daraus folgt die Behauptung über die Konvergenz der Partialbruch-Reihe des Cotangens.

Wir setzen

$$g(x) := \frac{1}{x} + \sum_{n=1}^{\infty} \frac{2x}{x^2 - n^2} = \lim_{N \rightarrow \infty} g_N(x) \quad \text{mit} \quad g_N(x) := \frac{1}{x} + \sum_{n=1}^N \frac{2x}{x^2 - n^2}.$$

Da $\frac{2x}{x^2 - n^2} = \frac{1}{x - n} + \frac{1}{x + n}$, gilt

$$g_N(x) = \sum_{n=-N}^N \frac{1}{x - n}.$$

Differenzieren ergibt für $x \in \mathbb{R} \setminus \mathbb{Z}$

$$g'_N(x) = \sum_{n=-N}^N \frac{-1}{(x - n)^2}.$$

Dies konvergiert für $N \rightarrow \infty$ auf jedem kompakten Intervall $I \subset \mathbb{R} \setminus \mathbb{Z}$ gleichmäßig gegen $\sum_{-\infty}^{\infty} \frac{-1}{(z - n)^2}$, also folgt aus Satz 5 und Beispiel (21.12)

$$g'(x) = -\left(\frac{\pi}{\sin \pi x}\right)^2 \quad \text{für alle } x \in \mathbb{R} \setminus \mathbb{Z}.$$

Andererseits gilt ebenfalls

$$\frac{d}{dx}(\pi \cot \pi x) = \frac{d}{dx}\left(\frac{\pi \cos \pi x}{\sin \pi x}\right) = -\left(\frac{\pi}{\sin \pi x}\right)^2.$$

Daraus folgt auf dem Intervall $0 < x < 1$

$$(*) \quad \pi \cot \pi x = g(x) + \text{const.}$$

Wir zeigen jetzt, dass die Konstante gleich 0 ist und die Gleichung für alle $x \in \mathbb{R} \setminus \mathbb{Z}$ gilt. Es ist

$$g(x+1) = \lim_{N \rightarrow \infty} g_N(x+1) = \lim_{N \rightarrow \infty} \sum_{n=-N}^N \frac{1}{x+1-n}$$

$$= \lim_{N \rightarrow \infty} \sum_{n=-N-1}^{N-1} \frac{1}{x-n} = \lim_{N \rightarrow \infty} \sum_{n=-N}^N \frac{1}{x-n} = g(x),$$

d.h. g ist (ebenso wie die Funktion $x \mapsto \cot \pi x$) periodisch mit der Periode 1. Außerdem ist g eine ungerade Funktion, d.h. $g(-x) = -g(x)$. Daraus folgt

$$g\left(\frac{1}{2}\right) = -g\left(-\frac{1}{2}\right) = -g\left(-\frac{1}{2} + 1\right) = -g\left(\frac{1}{2}\right) \implies g\left(\frac{1}{2}\right) = 0.$$

Da auch $\pi \cot \frac{\pi}{2} = 0$, gilt (*) mit $\text{const} = 0$ und wegen der Periodizität folgt $\pi \cot \pi x = g(x)$ für alle $x \in \mathbb{R} \setminus \mathbb{Z}$. Damit ist a) bewiesen.

b) Wir zeigen zunächst die Konvergenz des unendlichen Produkts. Für $M \geq N \geq 1$ sei

$$G_{NM}(x) := \prod_{n=N}^M \left(1 - \frac{x^2}{n^2}\right).$$

Sei $R > 0$ beliebig vorgegeben. Dann gilt für alle $x \in [-R, R]$ und $M \geq N \geq 2R$

$$\log G_{NM}(x) = \sum_{n=N}^M \log \left(1 - \frac{x^2}{n^2}\right).$$

Da $|x/n|^2 \leq R^2/n^2 \leq 1/4$ und $|\log(1-u)| \leq 2|u|$ für $|u| \leq 1/2$ (siehe Aufgabe 12.7), folgt

$$\left| \log \left(1 - \frac{x^2}{n^2}\right) \right| \leq \frac{2R^2}{n^2},$$

also konvergiert die Reihe $\sum_{n=N}^{\infty} \log(1 - x^2/n^2)$ nach dem Weierstraß'schen Majoranten-Kriterium (Satz 2) gleichmäßig auf dem Intervall $[-R, R]$. Durch Anwendung der Exponentialfunktion folgt daraus die gleichmäßige Konvergenz des unendlichen Produkts $\prod_{n=N}^{\infty} (1 - x^2/n^2)$, also auch des gesamten Produkts $\prod_{n=1}^{\infty}$. Siehe dazu Aufgabe 21.4.

Wir zeigen jetzt die Gleichung

$$(P) \quad \frac{\sin \pi x}{\pi x} = \prod_{n=1}^{\infty} \left(1 - \frac{x^2}{n^2}\right) \quad \text{für } |x| < 1,$$

wobei die linke Seite für $x = 0$ als 1 definiert wird (dadurch entsteht eine beliebig oft differenzierbare Funktion, wie die Potenzreihen-Darstellung zeigt, vgl. (21.11)).

Zum Beweis von (P) benutzen wir die logarithmische Ableitung

$$\frac{d}{dx} \log f(x) = \frac{f'(x)}{f(x)}.$$

Anwendung dieses Operators auf die linke Seite der Gleichung ergibt

$$\frac{d}{dx} \log \frac{\sin \pi x}{\pi x} = \pi \cot \pi x - \frac{1}{x} \quad \text{für } x \neq 0.$$

Für $x = 0$ ist der Wert der logarithmischen Ableitung der linken Seite gleich 0.

Anwendung auf einen Faktor der rechten Seite ergibt

$$\frac{d}{dx} \log \left(1 - \frac{x^2}{n^2} \right) = \frac{-2x/n^2}{1 - x^2/n^2} = \frac{2x}{x^2 - n^2},$$

also

$$\frac{d}{dx} \log \prod_{n=1}^{\infty} \left(1 - \frac{x^2}{n^2} \right) = \frac{d}{dx} \sum_{n=1}^{\infty} \log \left(1 - \frac{x^2}{n^2} \right) = \sum_{n=1}^{\infty} \frac{2x}{x^2 - n^2}.$$

Die Vertauschung von Differentiation und unendlicher Summe ist nach Satz 5 erlaubt.

Unter Verwendung von Teil a) erhält man daher im Intervall $|x| < 1$

$$\frac{d}{dx} \log(\text{Linke Seite}) = \frac{d}{dx} \log(\text{Rechte Seite}).$$

Daraus folgt, dass die linke Seite bis auf einen konstanten Faktor $\neq 0$ gleich der rechten Seite ist. Der Faktor muss aber gleich 1 sein, da für $x = 0$ beide Seiten den Wert 1 haben. Damit ist die Gleichung (P) bewiesen. Diese Gleichung gilt trivialerweise auch für $x = \pm 1$. Um die Produktformel des Sinus für alle $x \in \mathbb{R}$ zu zeigen, genügt es daher zu beweisen, dass die Funktion

$$G(x) := \lim_{N \rightarrow \infty} G_N(x), \quad \text{wobei} \quad G_N(x) := x \prod_{n=1}^N \left(1 - \frac{x^2}{n^2} \right),$$

die Periode 2 hat. Es gilt

$$G_N(x) = x \prod_{n=1}^N \frac{(n-x)(n+x)}{n^2} = \frac{(-1)^N}{N!^2} \prod_{n=-N}^N (x-n),$$

also

$$G_N(x+1) = \frac{(-1)^N}{N!^2} \prod_{n=-N-1}^{N-1} (x-n) = G_N(x) \frac{x+N+1}{x-N}.$$

Da $\lim_{N \rightarrow \infty} \frac{x+N+1}{x-N} = -1$, folgt $G(x+1) = -G(x)$, also $G(x+2) = G(x)$.

Damit ist Satz 7 vollständig bewiesen.

Wir geben zwei Anwendungen von Satz 7.

(21.13) Wir setzen in der Partialbruch-Zerlegung des Cotangens $x = 1/4$. Da $\cot(\pi/4) = 1$, folgt

$$\pi = 4 + \sum_{n=1}^{\infty} \left(\frac{1}{\frac{1}{4} - n} + \frac{1}{\frac{1}{4} + n} \right),$$

also

$$\frac{\pi}{4} = 1 + \sum_{n=1}^{\infty} \left(-\frac{1}{4n-1} + \frac{1}{4n+1} \right) = \sum_{k=0}^{\infty} \frac{(-1)^k}{2k+1}.$$

Damit haben wir die in (7.4) angegebene Formel für die Summe der Leibniz'schen Reihe bewiesen, vgl. auch (19.25). Ein weiterer Beweis wird in §22 mittels der Arcus-Tangens-Reihe gegeben.

(21.14) Setzt man im Sinus-Produkt $x = \frac{1}{2}$, erhält man

$$1 = \frac{\pi}{2} \prod_{n=1}^{\infty} \left(1 - \frac{1}{4n^2} \right) \implies \frac{\pi}{2} = \prod_{n=1}^{\infty} \frac{4n^2}{4n^2 - 1}.$$

Es ergibt sich also ein neuer Beweis für das Wallis'sche Produkt, siehe (19.24).

Corollar. Für alle $x \in \mathbb{R}$ mit $0 < x < 1$ gilt

$$\frac{\sin \pi x}{\pi} = \frac{1}{\Gamma(x)\Gamma(1-x)}.$$

Bemerkung. Die Einschränkung $0 < x < 1$ ist deshalb gemacht worden, damit beide Argumente der Gamma-Funktion im Nenner der rechten Seite positiv sind. Erweitert man jedoch den Definitions-Bereich der Gamma-Funktion gemäß Aufgabe 20.8 auf ganz $\mathbb{R}$ mit Ausnahme der ganzen Zahlen ≤ 0 , so gilt die Formel für alle $x \in \mathbb{R} \setminus \mathbb{Z}$.

Beweis. Wir verwenden die Gauß'sche Limesdarstellung der Gamma-Funktion (§20, Satz 5):

$$\frac{1}{\Gamma(x)} = \lim_{n \rightarrow \infty} \frac{x(x+1) \cdot \dots \cdot (x+n)}{n! n^x} = x \lim_{n \rightarrow \infty} n^{-x} \prod_{k=1}^n \left(1 + \frac{x}{k} \right)$$

und

$$\begin{aligned} \frac{1}{\Gamma(1-x)} &= \lim_{n \rightarrow \infty} \frac{(1-x)(2-x) \cdot \dots \cdot (n-x)(n+1-x)}{n! n^{1-x}} \\ &= \lim_{n \rightarrow \infty} n^x \frac{n+1-x}{n} \prod_{k=1}^n \left(1 - \frac{x}{k} \right). \end{aligned}$$

Multiplikation ergibt

$$\frac{1}{\Gamma(x)\Gamma(1-x)} = x \lim_{n \rightarrow \infty} \frac{n+1-x}{n} \prod_{k=1}^n \left(1 - \frac{x^2}{k^2}\right) = x \prod_{k=1}^{\infty} \left(1 - \frac{x^2}{k^2}\right).$$

Nach Satz 7 ist letzteres gleich $\frac{\sin \pi x}{\pi}$, q.e.d.

(21.15) Setzt man in der Formel des Corollars $x = \frac{1}{2}$, ergibt sich

$$\Gamma\left(\frac{1}{2}\right)^2 = \pi, \quad \text{d.h.} \quad \Gamma\left(\frac{1}{2}\right) = \sqrt{\pi},$$

was uns schon aus (20.10) bekannt ist.

AUFGABEN

21.1. Für $n \geq 1$ sei

$$f_n: \mathbb{R}_+ \rightarrow \mathbb{R}, \quad f_n(x) := \frac{x}{n^2} e^{-x/n}.$$

Man zeige, dass die Folge (f_n) auf $\mathbb{R}_+$ gleichmäßig gegen 0 konvergiert, aber

$$\lim_{n \rightarrow \infty} \int_0^{\infty} f_n(x) dx = 1.$$

21.2. Auf dem kompakten Intervall $[a, b] \subset \mathbb{R}$ seien $f_n: [a, b] \rightarrow \mathbb{R}$, $n \in \mathbb{N}$, Riemann-integrierbare Funktionen, die gleichmäßig gegen die Funktion $f: [a, b] \rightarrow \mathbb{R}$ konvergieren. Man zeige: Die Funktion f ist ebenfalls auf $[a, b]$ Riemann-integrierbar.

21.3. Sei $I =]a, b[\subset \mathbb{R}$ ein (eigentliches oder uneigentliches) Intervall ($a \in \mathbb{R} \cup \{-\infty\}$, $b \in \mathbb{R} \cup \{\infty\}$) und seien $f_n: I \rightarrow \mathbb{R}$, ($n \in \mathbb{N}$), stetige Funktionen, die auf jedem kompakten Teilintervall $[\alpha, \beta] \subset I$ gleichmäßig gegen die Funktion $f: I \rightarrow \mathbb{R}$ konvergieren. Es gebe eine nicht-negative Funktion $G: I \rightarrow \mathbb{R}$, die über I uneigentlich Riemann-integrierbar ist, so dass

$$|f_n(x)| \leq G(x) \quad \text{für alle } x \in I \text{ und } n \in \mathbb{N}.$$

Man zeige (Satz von der majorisierten Konvergenz):

Alle Funktionen f_n und f sind über I uneigentlich Riemann-integrierbar und

es gilt

$$\int_a^b f(x) dx = \lim_{n \rightarrow \infty} \int_a^b f_n(x) dx.$$

21.4. Seien $[a, b]$ und $[A, B]$ kompakte Intervalle in $\mathbb{R}$ und sei

$$f_n : [a, b] \longrightarrow [A, B] \subset \mathbb{R}, \quad n \in \mathbb{N},$$

eine Folge stetiger Funktionen, die gleichmäßig gegen eine Funktion $F : [a, b] \rightarrow \mathbb{R}$ konvergiert. Weiter sei $\varphi : [A, B] \rightarrow \mathbb{R}$ eine stetige Funktion. Man zeige: Die Folge der Funktionen

$$g_n := \varphi \circ f_n : [a, b] \rightarrow \mathbb{R}, \quad n \in \mathbb{N},$$

konvergiert gleichmäßig gegen die Funktion $G := \varphi \circ F$.

21.5. Für $|x| < 1$ berechne man die Summen der Reihen

$$\sum_{n=1}^{\infty} n^2 x^n, \quad \sum_{n=1}^{\infty} n^3 x^n \quad \text{und} \quad \sum_{n=1}^{\infty} \frac{x^n}{n}.$$

21.6. Man berechne die Summen der Reihen

$$\sum_{n=1}^{\infty} \frac{\sin nx}{n^3} \quad \text{und} \quad \sum_{n=1}^{\infty} \frac{\cos nx}{n^4}, \quad (x \in \mathbb{R}).$$

21.7. Sei $f(z) = \sum_0^{\infty} c_n (z-a)^n$ eine Potenzreihe mit komplexen Koeffizienten c_n . Sei R der Konvergenzradius dieser Reihe. Man zeige

$$R = \left(\limsup_{n \rightarrow \infty} \sqrt[n]{|c_n|} \right)^{-1} \quad (\text{Hadamardsche Formel}).$$

Dabei werde vereinbart $0^{-1} = \infty$ und $\infty^{-1} = 0$.

21.8. Man zeige, dass die Reihe

$$F(x) := \sum_{n=0}^{\infty} e^{-n^2 x}$$

für alle $x > 0$ konvergiert und eine beliebig oft differenzierbare Funktion $F : \mathbb{R}_+^* \rightarrow \mathbb{R}$ darstellt. Außerdem beweise man, dass für alle $k \geq 1$ gilt

$$\lim_{x \rightarrow \infty} F^{(k)}(x) = 0.$$

21.9. Sei $(a_n)_{n \geq 1}$ eine Folge reeller Zahlen. Die Reihe

$$f(x) = \sum_{n=1}^{\infty} \frac{a_n}{n^x}$$

konvergiere für ein $x_0 \in \mathbb{R}$. Man zeige: Die Reihe konvergiert gleichmäßig auf dem Intervall $[x_0, \infty[$.

21.10. Man beweise: Für alle $x \in \mathbb{R} \setminus \mathbb{Z}$ gilt

$$\frac{\pi}{\sin \pi x} = \frac{1}{x} + \sum_{n=1}^{\infty} (-1)^n \frac{2x}{x^2 - n^2} = \sum_{n \in \mathbb{Z}} \frac{(-1)^n}{x - n}.$$

21.11. Man beweise:

a) Die Produktdarstellung für $1/\Gamma(x)$ (Corollar zu Satz 5 aus §20)

$$\frac{1}{\Gamma(x)} = x e^{\gamma x} \prod_{n=1}^{\infty} \left(1 + \frac{x}{n}\right) e^{-x/n}, \quad (\gamma \text{ Euler-Mascheronische Konstante}),$$

konvergiert auf jedem Intervall $[\varepsilon, R]$, $0 < \varepsilon < R < \infty$, gleichmäßig.

b) Auf $\mathbb{R}_+^*$ gilt

$$-\log \Gamma(x) = \gamma x + \log x + \sum_{n=1}^{\infty} \left\{ \log \left(1 + \frac{x}{n}\right) - \frac{x}{n} \right\},$$

wobei die unendliche Reihe auf jedem Intervall $[\varepsilon, R]$, $0 < \varepsilon < R < \infty$, gleichmäßig konvergiert.

c) Die Gamma-Funktion ist auf $\mathbb{R}_+^*$ differenzierbar und es gilt

$$-\frac{\Gamma'(x)}{\Gamma(x)} = \gamma + \frac{1}{x} + \sum_{n=1}^{\infty} \left(\frac{1}{x+n} - \frac{1}{n} \right),$$

wobei die unendliche Reihe auf jedem Intervall $[\varepsilon, R]$, $0 < \varepsilon < R < \infty$, gleichmäßig konvergiert.

d) $\lim_{x \searrow 0} \left(\Gamma(x) - \frac{1}{x} \right) = \Gamma'(1) = -\gamma.$

§ 22 Taylor-Reihen

Wir haben schon die Darstellung verschiedener Funktionen, wie Exponentialfunktion, Sinus und Cosinus, durch Potenzreihen kennengelernt. In diesem Paragraphen beschäftigen wir uns systematisch mit der Entwicklung von Funktionen in Potenzreihen.

Als Erstes beweisen wir die Taylorsche Formel, die eine Approximation einer differenzierbaren Funktion durch ein Polynom mit einer Integraldarstellung des Fehlerterms gibt.

Hier und im ganzen Paragraphen sei $I \subset \mathbb{R}$ ein aus mehr als einem Punkt bestehendes Intervall.

Satz 1 (Taylorsche Formel). *Sei $f: I \rightarrow \mathbb{R}$ eine $(n+1)$ -mal stetig differenzierbare Funktion und $a \in I$. Dann gilt für alle $x \in I$*

$$f(x) = f(a) + \frac{f'(a)}{1!}(x-a) + \frac{f''(a)}{2!}(x-a)^2 + \dots + \frac{f^{(n)}(a)}{n!}(x-a)^n + R_{n+1}(x),$$

wobei

$$R_{n+1}(x) = \frac{1}{n!} \int_a^x (x-t)^n f^{(n+1)}(t) dt.$$

Beweis durch Induktion nach n .

Induktionsanfang. Für $n = 0$ ist die zu beweisende Formel

$$f(x) = f(a) + \int_a^x f'(t) dt$$

nichts anderes als der Fundamentalsatz der Differential- und Integralrechnung.

Induktionsschritt $n-1 \rightarrow n$. Nach Induktionsvoraussetzung ist

$$\begin{aligned} R_n(x) &= \frac{1}{(n-1)!} \int_a^x (x-t)^{n-1} f^{(n)}(t) dt \\ &= - \int_a^x f^{(n)}(t) \frac{d}{dt} \left(\frac{(x-t)^n}{n!} \right) dt = \quad (\text{Partielle Integration}) \end{aligned}$$

$$\begin{aligned}
&= -f^{(n)}(t) \frac{(x-t)^n}{n!} \Big|_{t=a}^{t=x} + \int_a^x \frac{(x-t)^n}{n!} df^{(n)}(t) \\
&= \frac{f^{(n)}(a)}{n!} (x-a)^n + \frac{1}{n!} \int_a^x (x-t)^n f^{(n+1)}(t) dt.
\end{aligned}$$

Daraus folgt die Behauptung.

Corollar. Sei $f: I \rightarrow \mathbb{R}$ eine $(n+1)$ -mal differenzierbare Funktion mit $f^{(n+1)}(x) = 0$ für alle $x \in I$. Dann ist f ein Polynom vom Grad $\leq n$.

Satz 2 (Lagrangesche Form des Restglieds). Sei $f: I \rightarrow \mathbb{R}$ eine $(n+1)$ -mal stetig differenzierbare Funktion und $a, x \in I$. Dann existiert ein ξ zwischen a und x , so dass

$$f(x) = \sum_{k=0}^n \frac{f^{(k)}(a)}{k!} (x-a)^k + \frac{f^{(n+1)}(\xi)}{(n+1)!} (x-a)^{n+1}.$$

Beweis. Nach dem Mittelwertsatz der Integralrechnung (§18, Satz 7) existiert ein $\xi \in [a, x]$ (bzw. $\xi \in [x, a]$, falls $x < a$), so dass gilt

$$\begin{aligned}
R_{n+1}(x) &= \frac{1}{n!} \int_a^x (x-t)^n f^{(n+1)}(t) dt = f^{(n+1)}(\xi) \int_a^x \frac{(x-t)^n}{n!} dt \\
&= -f^{(n+1)}(\xi) \frac{(x-t)^{n+1}}{(n+1)!} \Big|_a^x = \frac{f^{(n+1)}(\xi)}{(n+1)!} (x-a)^{n+1}, \quad \text{q.e.d.}
\end{aligned}$$

Corollar. Sei $f: I \rightarrow \mathbb{R}$ eine n -mal stetig differenzierbare Funktion und $a \in I$. Dann gilt für alle $x \in I$

$$f(x) = \sum_{k=0}^n \frac{f^{(k)}(a)}{k!} (x-a)^k + \varphi(x) (x-a)^n,$$

wobei φ eine Funktion mit $\lim_{x \rightarrow a} \varphi(x) = 0$ ist.

Beweis. Wir verwenden die Lagrangesche Form des Restglieds n -ter Ordnung

$$\begin{aligned}
f(x) - \sum_{k=0}^{n-1} \frac{f^{(k)}(a)}{k!} (x-a)^k &= \frac{f^{(n)}(\xi)}{n!} (x-a)^n \\
&= \frac{f^{(n)}(a)}{n!} (x-a)^n + \frac{f^{(n)}(\xi) - f^{(n)}(a)}{n!} (x-a)^n.
\end{aligned}$$

Wir setzen $\varphi(x) := \frac{f^{(n)}(\xi) - f^{(n)}(a)}{n!}$, (ξ hängt von x ab!).

Da ξ zwischen x und a liegt, folgt aus der Stetigkeit von $f^{(n)}$:

$$\lim_{x \rightarrow a} \varphi(x) = \lim_{\xi \rightarrow a} \frac{f^{(n)}(\xi) - f^{(n)}(a)}{n!} = 0.$$

Daraus folgt die Behauptung.

Bemerkung. Unter Verwendung des Landau-Symbols o lässt sich die Aussage des Corollars auch schreiben als

$$f(x) = \sum_{k=0}^n \frac{f^{(k)}(a)}{k!} (x-a)^k + o(|x-a|^n) \quad \text{für } x \rightarrow a.$$

Dies bedeutet, dass sich eine in einer Umgebung des Punktes a n -mal stetig differenzierbare Funktion f bis auf einen Fehler der Ordnung $o(|x-a|^n)$ durch das *Taylor-Polynom* n -ter Ordnung

$$T_n[f, a](x) := \sum_{k=0}^n \frac{f^{(k)}(a)}{k!} (x-a)^k$$

approximieren lässt.

(22.1) Als Beispiel betrachten wir die Funktion

$$f:]-1, 1[\rightarrow \mathbb{R}, \quad f(x) = \sqrt{1+x}$$

und den Entwicklungspunkt $a = 0$. Da

$$f(0) = 1, \quad f'(0) = \frac{1}{2\sqrt{1+x}} \Big|_{x=0} = \frac{1}{2},$$

ergibt das Corollar

$$f(x) = \sqrt{1+x} = 1 + \frac{x}{2} + \varphi(x)x \quad \text{mit } \lim_{x \rightarrow 0} \varphi(x) = 0.$$

Daraus erhält man z.B. für alle $n > 1$:

$$\begin{aligned} \sqrt{n+\sqrt{n}} &= \sqrt{n \left(1 + \frac{1}{\sqrt{n}} \right)} = \sqrt{n} \sqrt{1 + \frac{1}{\sqrt{n}}} \\ &= \sqrt{n} \left(1 + \frac{1}{2\sqrt{n}} + \varphi \left(\frac{1}{\sqrt{n}} \right) \frac{1}{\sqrt{n}} \right) = \sqrt{n} + \frac{1}{2} + \varphi \left(\frac{1}{\sqrt{n}} \right). \end{aligned}$$

Da $\lim_{n \rightarrow \infty} \varphi(1/\sqrt{n}) = 0$, folgt daraus

$$\lim_{n \rightarrow \infty} \left(\sqrt{n + \sqrt{n}} - \sqrt{n} \right) = \frac{1}{2}, \quad (\text{vgl. Aufgabe 6.7}).$$

Definition. Sei $f: I \rightarrow \mathbb{R}$ eine beliebig oft differenzierbare Funktion und $a \in I$. Dann heißt

$$T[f, a](x) := \sum_{k=0}^{\infty} \frac{f^{(k)}(a)}{k!} (x-a)^k$$

die *Taylor-Reihe* von f mit Entwicklungspunkt a .

Bemerkungen

- a) Der Konvergenzradius der Taylor-Reihe ist nicht notwendig > 0 .
- b) Falls die Taylor-Reihe von f konvergiert, konvergiert sie nicht notwendig gegen f .
- c) Die Taylor-Reihe konvergiert genau für diejenigen $x \in I$ gegen $f(x)$, für die das Restglied aus Satz 1 gegen 0 konvergiert.

(22.2) Wir geben ein Beispiel zu b). Sei $f: \mathbb{R} \rightarrow \mathbb{R}$ die Funktion (s. Bild 22.1)

$$f(x) := \begin{cases} e^{-1/x^2}, & \text{falls } x \neq 0, \\ 0, & \text{falls } x = 0. \end{cases}$$

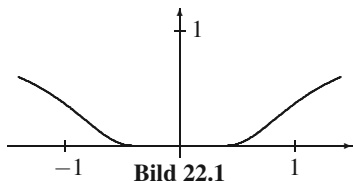


Bild 22.1

Dies ist eine Funktion, deren Graph sich in der Nähe des Nullpunkts sehr stark an die x -Achse anschmiegt (z.B. gilt $0 < f(x) < 10^{-10}$ für $0 < |x| \leq 0.2$).

Wir wollen zeigen, dass f beliebig oft differenzierbar ist und $f^{(n)}(0) = 0$ für alle $n \in \mathbb{N}$. Die Taylor-Reihe von f um den Nullpunkt ist also identisch 0, obwohl f selbst nur im Nullpunkt den Wert 0 annimmt.

Dazu beweisen wir durch vollständige Induktion nach n , dass es Polynome p_n gibt, so dass

$$f^{(n)}(x) = \begin{cases} p_n\left(\frac{1}{x}\right) e^{-1/x^2}, & \text{falls } x \neq 0, \\ 0, & \text{falls } x = 0. \end{cases}$$

Der Induktionsanfang $n = 0$ ist klar.

Induktionsschritt $n \rightarrow n + 1$.

a) Für $x \neq 0$ gilt

$$\begin{aligned} f^{(n+1)}(x) &= \frac{d}{dx} f^{(n)}(x) = \frac{d}{dx} \left(p_n \left(\frac{1}{x} \right) e^{-1/x^2} \right) \\ &= \left(-p'_n \left(\frac{1}{x} \right) \frac{1}{x^2} + 2p_n \left(\frac{1}{x} \right) \frac{1}{x^3} \right) e^{-1/x^2}. \end{aligned}$$

Man wähle $p_{n+1}(t) := -p'_n(t)t^2 + 2p_n(t)t^3$.

b) Für $x = 0$ gilt

$$\begin{aligned} f^{(n+1)}(0) &= \lim_{x \rightarrow 0} \frac{f^{(n)}(x) - f^{(n)}(0)}{x} = \lim_{x \rightarrow 0} \frac{p_n \left(\frac{1}{x} \right) e^{-1/x^2}}{x} \\ &= \lim_{R \rightarrow \pm\infty} R p_n(R) e^{-R^2} = 0 \text{ nach (12.1), q.e.d.} \end{aligned}$$

Aus §21, Corollar zu Satz 6 folgt unmittelbar

Satz 3. Sei $a \in \mathbb{R}$ und

$$f(x) = \sum_{n=0}^{\infty} c_n (x-a)^n$$

eine Potenzreihe mit einem positiven Konvergenzradius $r \in]0, \infty]$. Dann ist die Taylor-Reihe der Funktion $f:]a-r, a+r[\rightarrow \mathbb{R}$ mit Entwicklungspunkt a gleich dieser Potenzreihe (und konvergiert somit gegen f).

Beispiele

(22.3) Die Taylor-Reihe der Exponentialfunktion mit Entwicklungspunkt 0 ist

$$\exp(x) = \sum_{n=0}^{\infty} \frac{x^n}{n!}.$$

Sie konvergiert, wie wir bereits wissen, für alle $x \in \mathbb{R}$. Für einen beliebigen Entwicklungspunkt $a \in \mathbb{R}$ erhält man aus der Funktionalgleichung

$$\exp(x) = \exp(a) \exp(x-a) = \sum_{n=0}^{\infty} \frac{\exp(a)}{n!} (x-a)^n.$$

(22.4) Die Taylor-Reihen von Sinus und Cosinus,

$$\sin x = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k+1}}{(2k+1)!},$$

$$\cos x = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k}}{(2k)!},$$

konvergieren ebenfalls für alle $x \in \mathbb{R}$.

Die Lagrangesche Form des Restglieds ergibt für den Sinus

$$\sin x = \sum_{k=0}^n (-1)^k \frac{x^{2k+1}}{(2k+1)!} + R_{2n+3}(x)$$

mit

$$R_{2n+3}(x) = \frac{\sin^{(2n+3)}(\xi)}{(2n+3)!} x^{2n+3} = (-1)^{n+1} \frac{\cos \xi}{(2n+3)!} x^{2n+3}.$$

Dabei ist ξ eine Stelle zwischen 0 und x . Also gilt

$$|R_{2n+3}(x)| \leq \frac{|x|^{2n+3}}{(2n+3)!} \quad \text{für alle } x \in \mathbb{R}.$$

In §14, Satz 5, konnte diese Abschätzung nur für $|x| \leq 2n+4$ bewiesen werden. Ebenso beweist man für den Cosinus

$$\cos x = \sum_{k=0}^n (-1)^k \frac{x^{2k}}{(2k)!} + R_{2n+2}(x)$$

mit

$$R_{2n+2}(x) = (-1)^{n+1} \frac{\cos \xi}{(2n+2)!} x^{2n+2},$$

also

$$|R_{2n+2}(x)| \leq \frac{|x|^{2n+2}}{(2n+2)!} \quad \text{für alle } x \in \mathbb{R}.$$

Logarithmus und Arcus-Tangens

Wir bestimmen jetzt die Taylor-Reihen der Funktionen Logarithmus und Arcus-Tangens. Diese können beide durch Integration der Reihen ihrer Ableitungen gewonnen werden. Als spezielle Werte ergeben sich Formeln für die alternierende harmonische Reihe und die Leibniz'sche Reihe.

Satz 4 (Logarithmus-Reihe). Für $-1 < x \leq +1$ gilt

$$\log(1+x) = x - \frac{x^2}{2} + \frac{x^3}{3} \mp \dots = \sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{n} x^n.$$

Corollar. Für beliebiges $a > 0$ und $0 < x \leq 2a$ gilt

$$\log x = \log a + \sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{na^n} (x-a)^n.$$

Dies folgt aus der Funktionalgleichung, denn

$$\log x = \log(a + (x-a)) = \log a(1 + \frac{x-a}{a}) = \log a + \log(1 + \frac{x-a}{a}).$$

In Bild 22.2 sind die ersten drei Partialsummen der Taylor-Reihe des Logarithmus mit Entwicklungspunkt $a = 1$ dargestellt.

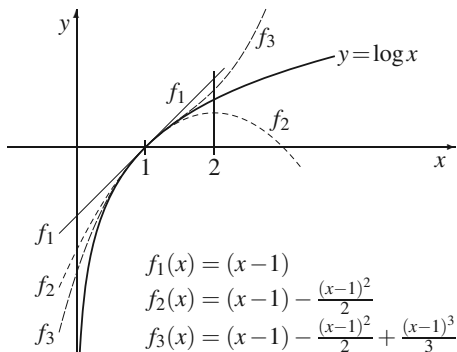


Bild 22.2 Taylor-Approximation des Logarithmus

Beweis von Satz 4. Für $|x| < 1$ gilt

$$\log(1+x) = \log(1+t) \Big|_0^x = \int_0^x \frac{dt}{1+t} = \int_0^x \left(\sum_{n=0}^{\infty} (-1)^n t^n \right) dt.$$

Nach §21, Satz 3, konvergiert $\sum_{n=0}^{\infty} (-1)^n t^n$ gleichmäßig auf $[-|x|, |x|]$. Aus §21, Satz 4, folgt daher

$$\log(1+x) = \sum_{n=0}^{\infty} (-1)^n \int_0^x t^n dt = \sum_{n=0}^{\infty} (-1)^n \frac{x^{n+1}}{n+1} = \sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{n} x^n.$$

Damit ist der Satz für $|x| < 1$ bewiesen. Um den noch fehlenden Fall $x = 1$ zu erledigen, beweisen wir zunächst ein allgemeines Resultat.

Satz 5 (Abelscher Grenzwertsatz). *Sei $\sum_{n=0}^{\infty} c_n$ eine konvergente Reihe reeller Zahlen. Dann konvergiert die Potenzreihe*

$$f(x) := \sum_{n=0}^{\infty} c_n x^n$$

gleichmäßig auf dem Intervall $[0, 1]$, stellt also dort eine stetige Funktion dar.

Bemerkung. Es gilt dann $\lim_{x \nearrow 1} \sum_{n=0}^{\infty} c_n x^n = \sum_{n=0}^{\infty} c_n$. Dies erklärt den Namen “Grenzwertsatz”.

Beweis. Nach §21, Satz 3, konvergiert die Reihe für alle x mit $|x| < 1$, also nach Voraussetzung auch für alle $x \in [0, 1]$. Es ist also nur zu beweisen, dass der Reihenrest

$$R_k(x) := \sum_{n=k}^{\infty} c_n x^n$$

für $k \rightarrow \infty$ auf $[0, 1]$ gleichmäßig gegen 0 konvergiert. Wir setzen

$$s_n := \sum_{k=n+1}^{\infty} c_k \quad \text{für } n \geq -1.$$

Es gilt $s_n - s_{n-1} = -c_n$ für alle $n \in \mathbb{N}$, und $\lim_{n \rightarrow \infty} s_n = 0$. Da die Folge der s_n beschränkt ist, konvergiert nach dem Majoranten-Kriterium die Reihe $\sum_{n=0}^{\infty} s_n x^n$ für $|x| < 1$. Nun ist

$$\begin{aligned} \sum_{n=k}^{\ell} c_n x^n &= - \sum_{n=k}^{\ell} s_n x^n + \sum_{n=k}^{\ell} s_{n-1} x^n \\ &= -s_{\ell} x^{\ell} - \sum_{n=k}^{\ell-1} s_n x^n + s_{k-1} x^k + \sum_{n=k}^{\ell-1} s_n x^{n+1} \\ &= -s_{\ell} x^{\ell} + s_{k-1} x^k - \sum_{n=k}^{\ell-1} s_n x^n (1-x). \end{aligned}$$

Der Grenzübergang $\ell \rightarrow \infty$ liefert für alle $x \in [0, 1]$

$$R_k(x) = s_{k-1} x^k - \sum_{n=k}^{\infty} s_n x^n (1-x).$$

Sei $\varepsilon > 0$ beliebig vorgegeben und $N \in \mathbb{N}$ so groß, dass $|s_n| < \varepsilon/2$ für alle $n \geq N$. Dann gilt für alle $k > N$ und alle $x \in [0, 1]$

$$|R_k(x)| < \frac{\varepsilon}{2} + \frac{\varepsilon}{2}(1-x) \sum_{n=k}^{\infty} x^n \leq \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon.$$

Damit ist Satz 5 bewiesen.

Nun zur Vervollständigung des Beweises von Satz 4!

Da $\sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{n}$ konvergiert, stellt nach dem Abelschen Grenzwertsatz

$$f(x) := \sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{n} x^n$$

eine in $[0, 1]$ stetige Funktion dar. Die Funktion $\log(1+x)$ ist ebenfalls in $[0, 1]$ stetig. Da $f(x) = \log(1+x)$ für alle $x \in [0, 1[$, gilt die Gleichung auch für $x = 1$. Damit haben wir die in (7.3) angegebene Formel

$$\log 2 = 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} \pm \dots$$

bewiesen. Diese Formel ist natürlich für die praktische Berechnung von $\log 2$ ganz ungeeignet. Will man hiermit $\log 2$ mit einer Genauigkeit von 10^{-k} berechnen, so muss man die ersten 10^k Glieder berücksichtigen. Zu einer besser konvergenten Reihe kommt man durch folgende Umformung: Für $|x| < 1$ gilt

$$\begin{aligned} \log(1+x) &= x - \frac{x^2}{2} + \frac{x^3}{3} - \frac{x^4}{4} \pm \dots, \\ \log(1-x) &= -x - \frac{x^2}{2} - \frac{x^3}{3} - \frac{x^4}{4} - \dots. \end{aligned}$$

Subtraktion ergibt

$$\log \frac{1+x}{1-x} = 2 \left(x + \frac{x^3}{3} + \frac{x^5}{5} + \dots \right) = 2 \sum_{k=0}^{\infty} \frac{x^{2k+1}}{2k+1}.$$

Dabei ist $\frac{1+x}{1-x} = y$, falls $x = \frac{y-1}{y+1}$. Für $x = \frac{1}{3}$ erhält man deshalb

$$\log 2 = 2 \left(\frac{1}{3} + \frac{1}{3 \cdot 3^3} + \frac{1}{5 \cdot 3^5} + \frac{1}{7 \cdot 3^7} + \dots \right).$$

Für eine Genauigkeit von 10^{-n} braucht man hier nur etwas mehr als n Glieder zu berücksichtigen; bricht man die Reihe mit dem Term $\frac{2}{(2k+1)3^{2k+1}}$ ab, so hat

man für den Fehler die Abschätzung

$$|R| \leq \frac{2}{(2k+3)3^{2k+3}} \sum_{m=0}^{\infty} \frac{1}{9^m} < \frac{1}{2k+3} \cdot \frac{1}{9^k}.$$

Z.B. erhält man mit $k = 48$ den $\log 2$ auf 45 Dezimalstellen genau:

$$\log 2 = 0.693147180559945309417232121458176568075500134 \dots$$

Zur Berechnung der Logarithmen anderer Argumente kann man die Funktionalgleichung heranziehen. Jede reelle Zahl $x > 0$ lässt sich schreiben als

$$x = 2^n \cdot y \quad \text{mit } n \in \mathbb{Z} \text{ und } 1 \leq y < 2.$$

Dann ist

$$\log x = n \log 2 + \log y = n \log 2 + \log \frac{1+z}{1-z},$$

wobei $z = \frac{y-1}{y+1}$, also insbesondere $|z| < \frac{1}{3}$.

Satz 6 (Arcus-Tangens-Reihe). Für $|x| \leq 1$ gilt

$$\arctan x = x - \frac{x^3}{3} + \frac{x^5}{5} - \frac{x^7}{7} \pm \dots = \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n+1}}{2n+1}.$$

Bemerkung. Man beachte, dass diese Potenzreihe bis auf die Vorzeichen bei den ungeraden Potenzen von x mit der Reihe für die Funktion $\frac{1}{2} \log \frac{1+x}{1-x}$ übereinstimmt. Dass dies kein reiner Zufall ist, darauf weist schon Aufgabe 14.4 hin. Der Zusammenhang wird klar in der Funktionentheorie, wo diese Funktionen auch für komplexe Argumente definiert werden. Dann gilt in der Tat die Formel $\arctan z = \frac{1}{2i} \log \frac{1+iz}{1-iz}$.

Beweis von Satz 6. Sei $|x| < 1$. Dann gilt

$$\begin{aligned} \arctan x &= \int_0^x \frac{dt}{1+t^2} = \int_0^x \left(\sum_{n=0}^{\infty} (-1)^n t^{2n} \right) dt \\ &= \sum_{n=0}^{\infty} (-1)^n \int_0^x t^{2n} dt = \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n+1}}{2n+1}. \end{aligned}$$

Dabei wurde §21, Satz 4, verwendet. Der Fall $|x| = 1$ wird analog zur Logarithmus-Reihe mithilfe des Abelschen Grenzwertsatzes bewiesen.

Da $\tan \frac{\pi}{4} = 1$, also $\arctan 1 = \frac{\pi}{4}$, ergibt sich für $x = 1$ die schon in (7.4) angegebene Summe für die Leibniz'sche Reihe

$$\frac{\pi}{4} = 1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} \pm \dots$$

Wie bei der alternierenden harmonischen Reihe für $\log(2)$ ist dies zwar eine interessante Formel, aber zur praktischen Berechnung von π ungeeignet. Eine effizientere Methode der Berechnung von π mittels des Arcus-Tangens liefert die *Machinsche Formel*

$$\frac{\pi}{4} = 4 \arctan \frac{1}{5} - \arctan \frac{1}{239},$$

siehe dazu Aufgabe 22.3. (Eine Fülle von weiteren Algorithmen zur Berechnung von π werden in dem Buch [AH] beschrieben.)

Binomische Reihe

Eine sehr interessante Reihe, die als Spezialfälle sowohl den binomischen Lehrsatz als auch die geometrische Reihe enthält, ist die binomische Reihe. Sie ergibt sich als Taylor-Reihe der allgemeinen Potenz $x \mapsto x^\alpha$ mit Entwicklungspunkt 1.

Satz 7 (Binomische Reihe). *Sei $\alpha \in \mathbb{R}$. Dann gilt für $|x| < 1$*

$$(1+x)^\alpha = \sum_{n=0}^{\infty} \binom{\alpha}{n} x^n.$$

Dabei ist $\binom{\alpha}{n} = \prod_{k=1}^n \frac{\alpha-k+1}{k}$.

Bemerkung. Für $\alpha \in \mathbb{N}$ bricht die Reihe ab, (denn in diesem Fall ist $\binom{\alpha}{n} = 0$ für $n > \alpha$) und die Formel folgt aus dem binomischen Lehrsatz (§1, Satz 5).

Beweis

a) Berechnung der Taylor-Reihe von $f(x) = (1+x)^\alpha$ mit Entwicklungspunkt 0:

$$f^{(k)}(x) = \alpha(\alpha-1) \cdot \dots \cdot (\alpha-k+1)(1+x)^{\alpha-k} = k! \binom{\alpha}{k} (1+x)^{\alpha-k}.$$

Da also $\frac{f^{(k)}(0)}{k!} = \binom{\alpha}{k}$, lautet die Taylor-Reihe von f

$$T[f, 0](x) = \sum_{k=0}^{\infty} \binom{\alpha}{k} x^k.$$

b) Wir zeigen, dass die Taylor-Reihe für $|x| < 1$ konvergiert. Dazu verwenden wir das Quotienten-Kriterium. Wir dürfen annehmen, dass $\alpha \notin \mathbb{N}$ und $x \neq 0$. Sei $a_n := \binom{\alpha}{n} x^n$. Dann gilt

$$\left| \frac{a_{n+1}}{a_n} \right| = \left| \frac{\binom{\alpha}{n+1} x^{n+1}}{\binom{\alpha}{n} x^n} \right| = |x| \cdot \left| \frac{\alpha - n}{n+1} \right|.$$

Da $\lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right| = |x| \lim_{n \rightarrow \infty} \left| \frac{\alpha - n}{n+1} \right| = |x| < 1$, existiert zu θ mit $|x| < \theta < 1$ ein n_0 , so dass

$$\left| \frac{a_{n+1}}{a_n} \right| \leq \theta \quad \text{für alle } n \geq n_0.$$

Also konvergiert die Taylor-Reihe für $|x| < 1$.

c) Wir beweisen jetzt, dass die Taylor-Reihe gegen f konvergiert. Es ist zu zeigen, dass das Restglied für $|x| < 1$ gegen 0 konvergiert. Es stellt sich heraus, dass man mit der Lagrangeschen Form des Restgliedes nicht weiterkommt. Wir verwenden deshalb die Integral-Darstellung. (Ein kürzerer Weg zum Beweis ist in Aufgabe 22.4 beschrieben. Es soll aber hier wenigstens ein Beispiel für die Anwendung der Integral-Form des Restgliedes vorgeführt werden.)

$$\begin{aligned} R_{n+1}(x) &= \frac{1}{n!} \int_0^x (x-t)^n f^{(n+1)}(t) dt \\ &= (n+1) \binom{\alpha}{n+1} \int_0^x (x-t)^n (1+t)^{\alpha-n-1} dt \end{aligned}$$

1. Fall: $0 \leq x < 1$.

Wir setzen $C := \max(1, (1+x)^\alpha)$. Dann gilt für $0 \leq t \leq x$

$$0 \leq (1+t)^{\alpha-n-1} \leq (1+t)^\alpha \leq C,$$

also

$$\begin{aligned} |R_{n+1}(x)| &= (n+1) \left| \binom{\alpha}{n+1} \right| \int_0^x (x-t)^n (1+t)^{\alpha-n-1} dt \\ &\leq (n+1) \left| \binom{\alpha}{n+1} \right| C \int_0^x (x-t)^n dt = C \left| \binom{\alpha}{n+1} \right| x^{n+1}. \end{aligned}$$

Weil nach b) die Reihe $\sum_{k=0}^{\infty} \binom{\alpha}{k} x^k$ für $|x| < 1$ konvergiert, folgt

$$\lim_{k \rightarrow \infty} \left| \binom{\alpha}{k} \right| x^k = 0, \quad \text{daher} \quad \lim_{n \rightarrow \infty} R_{n+1}(x) = 0.$$

2. Fall: $-1 < x < 0$. Hier gilt

$$\begin{aligned}
 |R_{n+1}(x)| &= (n+1) \left| \binom{\alpha}{n+1} \int_0^{|x|} (x+t)^n (1-t)^{\alpha-n-1} dt \right| \\
 &= \left| \alpha \binom{\alpha-1}{n} \right| \int_0^{|x|} (|x|-t)^n (1-t)^{\alpha-n-1} dt \\
 &\leq \left| \alpha \binom{\alpha-1}{n} \right| \int_0^{|x|} (|x|-t)^n (1-t)^{\alpha-n-1} dt \\
 &= \left| \alpha \binom{\alpha-1}{n} \right| |x|^n \int_0^{|x|} (1-t)^{\alpha-1} dt \\
 &\leq C \left| \binom{\alpha-1}{n} x^n \right| \quad \text{mit } C := |\alpha| \int_0^{|x|} (1-t)^{\alpha-1} dt.
 \end{aligned}$$

Da nach b) die Reihe $\sum_{n=0}^{\infty} \binom{\alpha-1}{n} x^n$ für $|x| < 1$ konvergiert, folgt

$$\lim_{n \rightarrow \infty} R_{n+1}(x) = 0, \quad \text{q.e.d.}$$

Beispiele

(22.5) Da $\binom{-1}{n} = (-1)^n$, ergibt sich für $\alpha = -1$ aus der binomischen Reihe

$$\frac{1}{1+x} = \sum_{n=0}^{\infty} (-1)^n x^n \quad \text{für } |x| < 1.$$

Die geometrische Reihe ist also ein Spezialfall der binomischen Reihe.

(22.6) Für $\alpha = \frac{1}{2}$ lauten die ersten Binomialkoeffizienten

$$\begin{aligned}
 \binom{\frac{1}{2}}{0} &= 1, \quad \binom{\frac{1}{2}}{1} = \frac{1}{2}, \quad \binom{\frac{1}{2}}{2} = \frac{\frac{1}{2}(-\frac{1}{2})}{1 \cdot 2} = -\frac{1}{8}, \\
 \binom{\frac{1}{2}}{3} &= \frac{\frac{1}{2}(-\frac{1}{2})(-\frac{3}{2})}{1 \cdot 2 \cdot 3} = \frac{1}{16}, \quad \binom{\frac{1}{2}}{4} = \binom{\frac{1}{2}}{3} \frac{-\frac{5}{2}}{4} = -\frac{5}{128}.
 \end{aligned}$$

Also gilt für $|x| < 1$:

$$\sqrt{1+x} = 1 + \frac{1}{2}x - \frac{1}{8}x^2 + \frac{1}{16}x^3 - \frac{5}{128}x^4 + \text{Glieder höherer Ordnung.}$$

Man kann dies zur näherungsweisen Berechnung von Wurzeln benützen; z.B. ist

$$\begin{aligned}
 \sqrt{10} &= \sqrt{9 \cdot \frac{10}{9}} = 3 \sqrt{1 + \frac{1}{9}} = 3 \left(1 + \frac{1}{2 \cdot 9} - \frac{1}{8 \cdot 9^2} + \dots \right) \\
 &= 3 + \frac{1}{6} - \frac{1}{8 \cdot 27} + \dots = 3.162 \dots
 \end{aligned}$$

(22.7) Für $\alpha = -\frac{1}{2}$ ist

$$\binom{-\frac{1}{2}}{0} = 1, \quad \binom{-\frac{1}{2}}{1} = -\frac{1}{2}, \quad \binom{-\frac{1}{2}}{2} = \frac{3}{8}, \quad \binom{-\frac{1}{2}}{3} = -\frac{5}{16}.$$

Daher gilt für $|x| < 1$:

$$\frac{1}{\sqrt{1+x}} = 1 - \frac{1}{2}x + \frac{3}{8}x^2 - \frac{5}{16}x^3 + \text{Glieder höherer Ordnung}.$$

Anwendung (Kinetische Energie eines relativistischen Teilchens).

Nach A. Einstein beträgt die Gesamtenergie eines Teilchens der Masse m

$$E = mc^2.$$

Dabei ist c die Lichtgeschwindigkeit. Die Masse ist jedoch von der Geschwindigkeit v des Teilchens abhängig; es gilt

$$m = \frac{m_0}{\sqrt{1 - (v/c)^2}}.$$

Hier ist m_0 die Ruhemasse des Teilchens; die Ruhenergie ist demnach $E_0 = m_0c^2$. Die kinetische Energie ist definiert als

$$E_{\text{kin}} = E - E_0.$$

Da $v < c$, kann man zur Berechnung die binomische Reihe verwenden:

$$\begin{aligned} E_{\text{kin}} &= mc^2 - m_0c^2 = m_0c^2 \left(\frac{1}{\sqrt{1 - (v/c)^2}} - 1 \right) \\ &= m_0c^2 \left(\frac{1}{2} \left(\frac{v}{c} \right)^2 + \frac{3}{8} \left(\frac{v}{c} \right)^4 + \dots \right) \\ &= \frac{1}{2}m_0v^2 + \frac{3}{8}m_0v^2 \left(\frac{v}{c} \right)^2 + \text{Glieder höherer Ordnung}. \end{aligned}$$

Der Term $\frac{1}{2}m_0v^2$ repräsentiert die kinetische Energie im klassischen Fall ($v \ll c$), der Term $\frac{3}{8}m_0v^2 \left(\frac{v}{c} \right)^2$ ist das Glied niedrigster Ordnung der Abweichung zwischen dem relativistischen und nicht-relativistischen Fall.

Wir wollen noch untersuchen, in welchen Fällen die binomische Reihe am Rande des Konvergenzintervalls konvergiert und beweisen dazu folgenden Hilssatz.

Hilssatz. Sei $\alpha \in \mathbb{R} \setminus \mathbb{N}$. Dann gibt es eine Konstante $c = c(\alpha) > 0$, so dass folgende asymptotische Beziehung besteht:

$$\left| \binom{\alpha}{n} \right| \sim \frac{c}{n^{1+\alpha}} \quad \text{für } n \rightarrow \infty.$$

Beweis

a) Sei zunächst $\alpha < 0$. Wir setzen $x := -\alpha$. Es ist

$$\left| \binom{-x}{n} \right| = \left| \frac{-x(-x-1) \cdot \dots \cdot (-x-n+1)}{n!} \right| = \frac{x(x+1) \cdot \dots \cdot (x+n)}{n!(x+n)}.$$

Daraus folgt unter Verwendung von §20, Satz 5

$$\lim_{n \rightarrow \infty} n^{1-x} \left| \binom{-x}{n} \right| = \lim_{n \rightarrow \infty} \frac{x(x+1) \cdot \dots \cdot (x+n)}{n! n^x} \cdot \frac{n}{n+x} = \frac{1}{\Gamma(x)}.$$

Die Behauptung gilt also mit der Konstanten $c(\alpha) = \frac{1}{\Gamma(-\alpha)}$.

b) Sei $k-1 < \alpha < k$ mit einer natürlichen Zahl $k \geq 1$. Dann ist auf $\alpha' = \alpha - k$ Teil a) anwendbar, d.h.

$$\lim_{n \rightarrow \infty} \left| \binom{\alpha-k}{n} \right| n^{1+\alpha-k} = \frac{1}{\Gamma(k-\alpha)},$$

also auch

$$\lim_{n \rightarrow \infty} \left| \binom{\alpha-k}{n-k} \right| n^{1+\alpha-k} = \frac{1}{\Gamma(k-\alpha)}.$$

Da $\binom{\alpha}{n} = \frac{\alpha(\alpha-1) \cdot \dots \cdot (\alpha-k+1)}{n(n-1) \cdot \dots \cdot (n-k+1)} \binom{\alpha-k}{n-k}$, folgt

$$\begin{aligned} \lim_{n \rightarrow \infty} \left| \binom{\alpha}{n} \right| n^{1+\alpha} &= \lim_{n \rightarrow \infty} \frac{n^k \alpha(\alpha-1) \cdot \dots \cdot (\alpha-k+1)}{n(n-1) \cdot \dots \cdot (n-k+1)} \left| \binom{\alpha-k}{n-k} \right| n^{1+\alpha-k} \\ &= \frac{\alpha(\alpha-1) \cdot \dots \cdot (\alpha-k+1)}{\Gamma(k-\alpha)} =: c(\alpha), \quad \text{q.e.d.} \end{aligned}$$

Zusatz zu Satz 7.

a) Für $\alpha \geq 0$ konvergiert die binomische Reihe

$$(1+x)^\alpha = \sum_{n=0}^{\infty} \binom{\alpha}{n} x^n$$

absolut und gleichmäßig im Intervall $[-1, +1]$.

b) Für $-1 < \alpha < 0$ konvergiert die binomische Reihe für $x = +1$ und divergiert für $x = -1$.

c) Für $\alpha \leq -1$ divergiert die binomische Reihe sowohl für $x = +1$ als auch für $x = -1$.

Beweis

a) Wir können annehmen, dass $\alpha \notin \mathbb{N}$, da für $\alpha \in \mathbb{N}$ die binomische Reihe abbricht. Aus dem Hilfssatz folgt, dass es eine Konstante $K > 0$ gibt mit

$$\left| \binom{\alpha}{n} \right| \leq \frac{K}{n^{1+\alpha}} \quad \text{für alle } n \geq 1.$$

Da die Reihe $\sum \frac{1}{n^{1+\alpha}}$ für $\alpha > 0$ konvergiert, folgt die Behauptung.

b) Für $-1 < \alpha < 0$ gilt $\binom{\alpha}{n} = (-1)^n \left| \binom{\alpha}{n} \right|$. Die Konvergenz der binomischen Reihe für $x = 1$ folgt nun aus dem Leibniz'schen Konvergenzkriterium für alternierende Reihen, die Divergenz an der Stelle $x = -1$ daraus, dass $\sum \frac{1}{n^{1+\alpha}}$ für $\alpha < 0$ divergiert.

c) Aus dem Hilfssatz folgt, dass $\binom{\alpha}{n}$ für $n \rightarrow \infty$ nicht gegen 0 konvergiert, falls $\alpha \leq -1$. Deshalb divergieren in diesem Fall die Reihen

$$\sum_{n=0}^{\infty} \binom{\alpha}{n} \quad \text{und} \quad \sum_{n=0}^{\infty} \binom{\alpha}{n} (-1)^n.$$

(22.8) Beispiel. Für $|x| \leq 1$ gilt auch $|x^2 - 1| \leq 1$. Also haben wir die im Intervall $[-1, 1]$ gleichmäßig konvergente Entwicklung

$$|x| = \sqrt{x^2} = \sqrt{1 + (x^2 - 1)} = \sum_{n=0}^{\infty} \binom{\frac{1}{2}}{n} (x^2 - 1)^n.$$

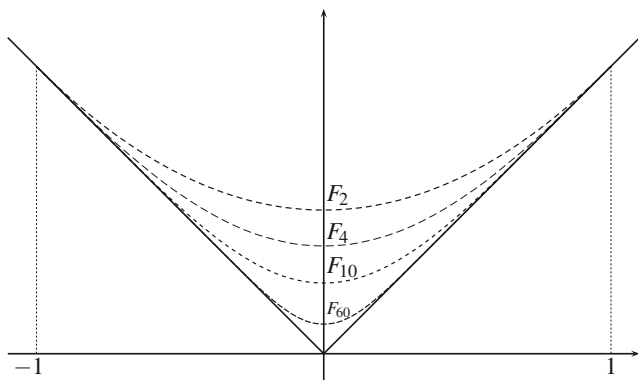
Die Funktion abs kann also in $[-1, 1]$ gleichmäßig durch Polynome approximiert werden. Wir wollen uns das etwas genauer anschauen. Wir bezeichnen die Partialsummen mit

$$F_{2n}(x) := \sum_{k=0}^n \binom{\frac{1}{2}}{k} (x^2 - 1)^k.$$

F_{2n} sind Polynome vom Grad $2n$, die für $n \rightarrow \infty$ auf dem Intervall $[-1, 1]$ gleichmäßig gegen die Funktion $x \mapsto |x|$ konvergieren. Die Konvergenz ist jedoch in der Nähe des Nullpunkts sehr langsam, siehe Bild 22.3.

Die ersten Polynome lauten ausgeschrieben

$$\begin{aligned} F_2(x) &= \frac{1}{2}(1 + x^2), \\ F_4(x) &= \frac{1}{8}(3 + 6x^2 - x^4), \end{aligned}$$

Bild 22.3 Approximation von $\text{abs}(x)$

$$F_6(x) = \frac{1}{16}(5 + 15x^2 - 5x^4 + x^6),$$

$$F_8(x) = \frac{1}{128}(35 + 140x^2 - 70x^4 + 28x^6 - 5x^8),$$

$$F_{10}(x) = \frac{1}{256}(63 + 315x^2 - 210x^4 + 126x^6 - 45x^8 + 7x^{10}).$$

Man beachte folgenden Unterschied zu den Taylor-Reihen: Während bei der Folge der Partialsummen der Taylor-Reihen immer nur Terme höherer Ordnung hinzukommen, ändern sich hier bei jedem Schritt die Koeffizienten aller auftretenden Potenzen von x .

Als Folgerung der Approximation der Funktion abs durch Polynome können wir nun den Weierstraß'schen Approximationssatz herleiten.

Satz 8 (Weierstraß'scher Approximationssatz). *Jede auf einem kompakten Intervall $[a, b] \subset \mathbb{R}$ stetige Funktion $f : [a, b] \rightarrow \mathbb{R}$ lässt sich gleichmäßig durch Polynome approximieren, d.h. zu jedem $\varepsilon > 0$ existiert ein Polynom P mit $\|f - P\| < \varepsilon$, wobei $\|$ die Supremumsnorm auf $[a, b]$ bezeichne.*

Beweis. Nach Aufgabe 11.8 existiert eine stückweise lineare stetige Funktion $\varphi \in \text{PL}[a, b]$ mit $\|f - \varphi\| < \varepsilon/2$. Die Funktion φ lässt sich schreiben als

$$\varphi(x) = \alpha + \beta x + \sum_{k=1}^r c_k |x - t_k|$$

mit $t_k \in]a, b[$ und Konstanten $\alpha, \beta, c_k \in \mathbb{R}$. Aus (22.8) folgt, dass es zu jeder Funktion $\varphi_k(x) := c_k|x - t_k|$ ein Polynom P_k gibt mit $\|\varphi_k - P_k\| < \varepsilon/(2r)$. Für das Polynom

$$P(x) := \alpha + \beta x + \sum_{k=1}^r P_k(x)$$

gilt dann $\|f - P\| < \varepsilon$, q.e.d.

Produkte und Quotienten von Potenzreihen

In diesem Abschnitt beweisen wir, dass Produkte und Quotienten (falls der Nenner ungleich null) von Funktionen, die sich durch Potenzreihen darstellen lassen, sich wieder in eine Potenzreihe entwickeln lassen. Unter anderem gelangen wir so zur Taylor-Reihe der Funktion Tangens.

Satz 9. Seien $f(z) = \sum_{n=0}^{\infty} a_n z^n$ und $g(z) = \sum_{n=0}^{\infty} b_n z^n$ Potenzreihen mit komplexen Koeffizienten und positivem Konvergenzradius r_1 bzw. r_2 . Dann wird das Produkt fg im Kreis $\{z \in \mathbb{C} : |z| < \min(r_1, r_2)\}$, dargestellt durch die Potenzreihe

$$f(z)g(z) = \sum_{n=0}^{\infty} c_n z^n \quad \text{mit} \quad c_n := \sum_{k=0}^n a_k b_{n-k}.$$

Beweis. Dies folgt unmittelbar aus dem Satz über das Cauchy-Produkt von unendlichen Reihen.

(22.9) Als Beispiel betrachten wir für $\alpha, \beta \in \mathbb{R}$ die binomischen Reihen

$$f(x) = (1+x)^\alpha = \sum_{n=0}^{\infty} \binom{\alpha}{n} x^n \quad \text{und} \quad g(x) = (1+x)^\beta = \sum_{n=0}^{\infty} \binom{\beta}{n} x^n$$

Ihr Cauchy-Produkt $f(x)g(x) = \sum_{n=0}^{\infty} c_n x^n$ hat die Koeffizienten

$$c_n = \sum_{k=0}^n \binom{\alpha}{k} \binom{\beta}{n-k}.$$

Andrerseits gilt $f(x)g(x) = (1+x)^{\alpha+\beta} = \sum_{n=0}^{\infty} \binom{\alpha+\beta}{n} x^n$. Aus dem Identitätssatz (§21, Corollar zu Satz 3) folgt die Formel

$$\binom{\alpha+\beta}{n} = \sum_{k=0}^n \binom{\alpha}{k} \binom{\beta}{n-k},$$

vgl. Aufgabe 1.11.

Da ein Quotient f/g sich als $f \cdot (1/g)$ schreiben lässt, genügt es, statt allgemeiner Quotienten von Potenzreihen nur das Reziproke einer Potenzreihe zu untersuchen.

Satz 10. Sei $f(z) = \sum_{n=0}^{\infty} a_n z^n$ eine Potenzreihe mit komplexen Koeffizienten und Konvergenzradius $0 < r \leq \infty$. Es gelte $f(0) = a_0 \neq 0$. Dann gibt es ein ρ mit $0 < \rho \leq r$, so dass $f(z) \neq 0$ für alle $z \in \mathbb{C}$ mit $|z| < \rho$ und sich $1/f$ im Kreis $\{z \in \mathbb{C} : |z| < \rho\}$ in eine Potenzreihe

$$\frac{1}{f(z)} = \sum_{n=0}^{\infty} b_n z^n$$

entwickeln lässt.

Beweis. Es genügt zu zeigen: Es gibt eine Potenzreihe $g(z) = \sum_{n=0}^{\infty} b_n z^n$ mit einem positiven Konvergenzradius $\rho \leq r$, so dass

$$f(z)g(z) = 1 \quad \text{für alle } |z| < \rho.$$

Ohne uns zunächst um die Konvergenz zu kümmern, untersuchen wir zuerst, welche Bedingungen die Koeffizienten b_n erfüllen müssen, dass $fg = 1$. Dazu muss das Cauchy-Produkt der Potenzreihen f und g gleich der trivialen Potenzreihe 1 sein. Dies bedeutet

$$(*) \left\{ \begin{array}{l} a_0 b_0 = 1, \\ a_0 b_1 + a_1 b_0 = 0, \\ a_0 b_2 + a_1 b_1 + a_2 b_0 = 0, \\ \vdots \\ a_0 b_n + a_1 b_{n-1} + \dots + a_n b_0 = 0, \\ \vdots \end{array} \right.$$

Daraus kann man, ausgehend von $b_0 := 1/a_0$, rekursiv alle b_n berechnen gemäß

$$b_n := -\frac{1}{a_0} \sum_{k=0}^{n-1} a_{n-k} b_k \quad \text{für } n \geq 1.$$

Es bleibt nur noch zu zeigen, dass die damit eindeutig bestimmte Potenzreihe einen positiven Konvergenzradius hat. Dies beweisen wir jetzt.

O.B.d.A. sei $a_0 = 1$ (andernfalls multipliziere man f mit der Konstanten $1/a_0$

und g mit a_0). Da die Reihe

$$f_1(z) := \sum_{n=1}^{\infty} a_n z^n$$

für $|z| < r$ absolut konvergiert und $f_1(0) = 0$, gibt es ein $0 < \rho < r$ mit

$$\sum_{n=1}^{\infty} |a_n| \rho^n \leq 1.$$

Behauptung. Es gilt $|b_n| \rho^n \leq 1$ für alle $n \geq 0$.

Wir zeigen die Behauptung durch vollständige Induktion. Der Induktionsanfang $n = 0$ ist trivial, da $b_0 = 1$.

Induktionsschritt. Sei schon bekannt, dass $|b_k| \rho^k \leq 1$ für alle $k < n$. Aus der Definition von b_n folgt

$$|b_n| \rho^n \leq \sum_{k=0}^{n-1} |a_{n-k}| \rho^{n-k} |b_k| \rho^k \leq \sum_{k=0}^{n-1} |a_{n-k}| \rho^{n-k} \leq \sum_{m=1}^{\infty} |a_m| \rho^m \leq 1.$$

Damit ist die Behauptung bewiesen. Es folgt, dass die Potenzreihe

$g(z) = \sum_0^{\infty} b_n z^n$ einen Konvergenzradius $\geq \rho$ hat, q.e.d.

(22.10) Bernoulli-Zahlen. Als Beispiel für Satz 10 betrachten wir die Funktion $f(z) := (e^z - 1)/z$, $f(0) := 1$. Aus der Exponentialreihe erhält man die Potenzreihen-Entwicklung

$$f(z) = \frac{e^z - 1}{z} = \sum_{n=0}^{\infty} \frac{z^n}{(n+1)!}$$

mit Konvergenzradius ∞ . Nach Satz 10 lässt sich $1/f(z)$ in einer gewissen Kreisscheibe $\{|z| < \rho\}$, $\rho > 0$, in eine Potenzreihe entwickeln. Die Koeffizienten dieser Potenzreihe sind bis auf einen Faktor $1/n!$ die sog. Bernoulli-Zahlen B_n . Es gilt also nach Definition

$$\frac{z}{e^z - 1} = \sum_{n=0}^{\infty} \frac{B_n}{n!} z^n.$$

Die B_n können berechnet werden durch Betrachtung des Cauchy-Produkts

$$\left(\sum_{k=0}^{\infty} \frac{B_k}{k!} z^k \right) \left(\sum_{\ell=0}^{\infty} \frac{1}{(\ell+1)!} z^\ell \right) = 1.$$

Dies führt auf die Gleichungen $B_0 = 1$ und

$$\sum_{k+\ell=n} \frac{B_k}{k!(\ell+1)!} = 0 \quad \text{für } n \geq 1.$$

Nach Multiplikation mit $(n+1)!$ ist letzteres äquivalent zu

$$\sum_{k=0}^n \binom{n+1}{k} B_k = 0, \quad (n \geq 1).$$

Die ersten dieser Gleichungen lauten ausgeschrieben

$$\begin{aligned} B_0 + 2B_1 &= 0 &\implies B_1 &= -\frac{1}{2}, \\ B_0 + 3B_1 + 3B_2 &= 0 &\implies B_2 &= \frac{1}{6}, \\ B_0 + 4B_1 + 6B_2 + 4B_3 &= 0 &\implies B_3 &= 0, \\ B_0 + 5B_1 + 10B_2 + 10B_3 + 5B_4 &= 0 &\implies B_4 &= -\frac{1}{30}. \end{aligned}$$

So fortfahrend kann man der Reihe nach alle Bernoulli-Zahlen, die sämtlich rational sind, berechnen. Wir geben eine Liste der nicht-verschwindenden B_n mit $n \leq 16$.

n	0	1	2	4	6	8	10	12	14	16
B_n	1	$-\frac{1}{2}$	$\frac{1}{6}$	$-\frac{1}{30}$	$\frac{1}{42}$	$-\frac{1}{30}$	$\frac{5}{66}$	$-\frac{691}{2730}$	$\frac{7}{6}$	$-\frac{3617}{510}$

Die Bernoulli-Zahlen haben interessante zahlentheoretische Eigenschaften. Sei $B_n = a_n/b_n$ mit teilerfremden ganzen Zahlen a_n, b_n . Beispielsweise liest man aus der Tabelle Folgendes ab: Ist $2k+1 = p$ eine Primzahl, so ist p ein Teiler von b_{2k} , (in Zeichen $p \mid b_{2k}$):

$$\begin{aligned} 3 \mid b_2 = 6, & \quad 5 \mid b_4 = 30, & \quad 7 \mid b_6 = 42, \\ 11 \mid b_{10} = 66, & \quad 13 \mid b_{12} = 2730, & \quad 17 \mid b_{16} = 510. \end{aligned}$$

Dies ist tatsächlich allgemein wahr. Genauer gilt folgender Satz von v. Staudt: Der Nenner b_{2k} ist gleich dem Produkt aller Primzahlen p , so dass $p-1$ ein Teiler von $2k$ ist. Siehe dazu [HWr], Chap. VII.

Außer B_1 verschwinden alle Bernoulli-Zahlen mit ungeradem Index. Dies kann man so beweisen:

$$\frac{z}{e^z - 1} + \frac{z}{2} = \frac{z}{2} \cdot \frac{e^z + 1}{e^z - 1} = \frac{z}{2} \cdot \frac{e^{z/2} + e^{-z/2}}{e^{z/2} - e^{-z/2}}.$$

Das ist eine gerade Funktion, d.h. invariant gegenüber der Substitution $z \mapsto -z$. Daher müssen in der Potenzreihen-Entwicklung alle Glieder ungerader

Ordnung verschwinden:

$$\frac{z}{e^z - 1} + \frac{z}{2} = \frac{z}{2} \cdot \frac{e^{z/2} + e^{-z/2}}{e^{z/2} - e^{-z/2}} = \sum_{k=0}^{\infty} \frac{B_{2k}}{(2k)!} z^{2k}.$$

Daraus ergibt sich $B_1 = -\frac{1}{2}$ und $B_{2k+1} = 0$ für $k \geq 1$.

(22.11) Potenzreihen-Entwicklung des Cotangens. Setzt man in der letzten Formel von (22.10) $z = 2ix$, so erhält man

$$\sum_{k=0}^{\infty} (-1)^k \frac{2^{2k} B_{2k}}{(2k)!} x^{2k} = ix \frac{e^{ix} + e^{-ix}}{e^{ix} - e^{-ix}} = x \frac{\cos x}{\sin x} = x \cot x.$$

Diese Gleichung gilt für $|x| < \rho/2$, wobei ρ den Konvergenzradius der Reihe $\sum_0^{\infty} (B_n/n!) z^n$ bezeichnet. (Wir werden anschließend zeigen, dass $\rho = 2\pi$.)

Substituiert man darin x durch πx und dividiert die Gleichung durch x , ergibt sich für $0 < |x| < \rho/2\pi$

$$\pi \cot \pi x = \frac{1}{x} + \sum_{k=1}^{\infty} (-1)^k \frac{(2\pi)^{2k} B_{2k}}{(2k)!} x^{2k-1}.$$

Andererseits gilt nach §21, Satz 7 für $x \in \mathbb{R} \setminus \mathbb{Z}$

$$\pi \cot \pi x = \frac{1}{x} + \sum_{n=1}^{\infty} \left(\frac{1}{x-n} + \frac{1}{x+n} \right).$$

Durch Vergleich erhält man für $|x| < \rho/2\pi$

$$f(x) := \sum_{k=1}^{\infty} (-1)^k \frac{(2\pi)^{2k} B_{2k}}{(2k)!} x^{2k-1} = \sum_{n=1}^{\infty} \left(\frac{1}{x-n} + \frac{1}{x+n} \right).$$

(Für $x = 0$ gilt die Gleichung trivialerweise.) Nach §21, Satz 5 kann man die rechte Seite gliedweise differenzieren und erhält

$$f^{(2k-1)}(x) = -(2k-1)! \sum_{n=1}^{\infty} \left(\frac{1}{(x-n)^{2k}} + \frac{1}{(x+n)^{2k}} \right),$$

also insbesondere

$$\frac{f^{(2k-1)}(0)}{(2k-1)!} = -2 \sum_{n=1}^{\infty} \frac{1}{n^{2k}} = -2\zeta(2k).$$

Aus der Potenzreihen-Entwicklung von f liest man ab (§21, Satz 6, Corollar):

$$\frac{f^{(2k-1)}(0)}{(2k-1)!} = (-1)^k \frac{(2\pi)^{2k} B_{2k}}{(2k)!}.$$

Fasst man die beiden Gleichungen zusammen, erhält man

Satz 11 (Euler). Für die Werte der Zetafunktion $\zeta(s) = \sum_{n=1}^{\infty} \frac{1}{n^s}$ an den positiven geraden Zahlen gilt:

$$\zeta(2k) = (-1)^{k-1} \frac{2^{2k-1} B_{2k}}{(2k)!} \pi^{2k} \quad \text{für alle } k \geq 1,$$

insbesondere

$$\zeta(2) = \frac{\pi^2}{6}, \quad \zeta(4) = \frac{\pi^4}{90}, \quad \zeta(6) = \frac{\pi^6}{945}, \quad \zeta(8) = \frac{\pi^8}{9450}, \quad \zeta(10) = \frac{\pi^{10}}{93555}.$$

Corollar. a) Für alle $k \geq 1$ ist $(-1)^{k-1} B_{2k} > 0$.

b) Es gilt $\lim_{k \rightarrow \infty} \sqrt[2k]{\frac{|B_{2k}|}{(2k)!}} = \frac{1}{2\pi}.$

c) Die Potenzreihe

$$\frac{z}{e^z - 1} = \sum_{n=0}^{\infty} \frac{B_n}{n!} z^n = 1 - \frac{z}{2} + \sum_{k=1}^{\infty} \frac{B_{2k}}{(2k)!} z^{2k}$$

hat den Konvergenzradius 2π .

Beweis. a) folgt daraus, dass $\zeta(2k) > 0$.

b) Es gilt $\frac{|B_{2k}|}{(2k)!} = \frac{2\zeta(2k)}{(2\pi)^{2k}}$ und $1 < \zeta(2k) < 2$, also $\lim_{k \rightarrow \infty} \sqrt[2k]{2\zeta(2k)} = 1$.

c) folgt unmittelbar aus b), vgl. Aufgabe 21.7.

Satz 12. a) Für $0 < |x| < \pi$ gilt

$$\begin{aligned} \cot x &= \frac{1}{x} + \sum_{k=1}^{\infty} (-1)^k \frac{2^{2k} B_{2k}}{(2k)!} x^{2k-1} \\ &= \frac{1}{x} - \frac{1}{3}x - \frac{1}{45}x^3 - \frac{2}{945}x^5 - \frac{1}{4725}x^7 - \dots \end{aligned}$$

b) Für $|x| < \pi/2$ gilt

$$\begin{aligned} \tan x &= \sum_{k=1}^{\infty} (-1)^{k-1} \frac{2^{2k} (2^{2k} - 1) B_{2k}}{(2k)!} x^{2k-1} \\ &= x + \frac{1}{3}x^3 + \frac{2}{15}x^5 + \frac{17}{315}x^7 + \dots \end{aligned}$$

Beweis. a) Dies folgt aus (22.11) und Teil c) des Corollars.

b) Wir benutzen die einfach zu beweisende Formel (siehe Aufgabe 14.8)

$$\tan x = \cot x - 2 \cot 2x \quad \text{für alle } x \in \mathbb{R} \setminus \mathbb{Z} \frac{\pi}{2}.$$

Setzt man darin die Entwicklung für den Cotangens aus a) ein, so heben sich die Glieder $\frac{1}{x}$ weg, und man erhält die Potenzreihen-Entwicklung des Tangens für $0 < |x| < \pi/2$. Für $x = 0$ gilt sie aber trivialerweise.

AUFGABEN

22.1. Sei $f :]a, b[\rightarrow \mathbb{R}$ eine n -mal stetig differenzierbare Funktion ($n \geq 1$). Im Punkt $x_0 \in]a, b[$ gelte:

$$f^{(k)}(x_0) = 0 \quad \text{für } 1 \leq k < n \quad \text{und} \quad f^{(n)}(x_0) \neq 0.$$

Man beweise mithilfe des Corollars zu Satz 2:

- a) Ist n ungerade, so besitzt f in x_0 kein lokales Extremum.
- b) Ist n gerade, so besitzt f in x_0 ein strenges lokales Maximum bzw. Minimum, je nachdem, ob $f^{(n)}(x_0) < 0$ oder $f^{(n)}(x_0) > 0$.

22.2. Sei $f :]a - \varepsilon, a + \varepsilon[\rightarrow \mathbb{R}$ eine n -mal stetig differenzierbare Funktion in einer Umgebung des Punktes $a \in \mathbb{R}$, ($\varepsilon > 0$). Es gelte

$$f(x) = c_0 + c_1(x-a) + c_2(x-a)^2 + \dots + c_n(x-a)^n + o(|x-a|^n)$$

Man zeige: Dann ist notwendig

$$c_k = \frac{1}{k!} f^{(k)}(a) \quad \text{für } k = 0, 1, \dots, n.$$

22.3. Man beweise die Funktionalgleichung der Arcus-Tangens:

Für $x, y \in \mathbb{R}$ mit $|\arctan x + \arctan y| < \frac{\pi}{2}$ gilt

$$\arctan x + \arctan y = \arctan \frac{x+y}{1-xy}.$$

Man folgere hieraus die "Machinsche Formel"

$$\frac{\pi}{4} = 4 \arctan \frac{1}{5} - \arctan \frac{1}{239}$$

und die Reihenentwicklung

$$\frac{\pi}{4} = \frac{4}{5} \sum_{k=0}^{\infty} \frac{(-1)^k}{2k+1} \left(\frac{1}{5}\right)^{2k} - \frac{1}{239} \sum_{k=0}^{\infty} \frac{(-1)^k}{2k+1} \left(\frac{1}{239}\right)^{2k}.$$

Welche Glieder muss man berücksichtigen, um π auf 1000 Dezimalstellen genau zu berechnen?

22.4. Diese Aufgabe beschreibt einen anderen Weg zum Beweis von Satz 7 über die binomische Reihe.

Man betrachte die auf dem Intervall $] -1, 1[$ definierte Funktion

$$f(x) := \sum_{n=0}^{\infty} \binom{\alpha}{n} x^n$$

und beweise für sie die Differentialgleichung

$$f'(x) = \frac{\alpha}{1+x} f(x).$$

Daraus leite man ab, dass die Funktion $g(x) := f(x)(1+x)^{-\alpha}$ konstant gleich 1 ist, also $f(x) = (1+x)^\alpha$ für $|x| < 1$ gilt.

22.5. Für einen reellen Parameter k mit $|k| < 1$ heißt

$$E(k) := \int_0^{\pi/2} \frac{dt}{\sqrt{1-k^2 \sin^2 t}}$$

vollständiges elliptisches Integral 1. Gattung.

Man entwickle $E(k)$ als Funktion von k in eine Taylor-Reihe, indem man $\frac{1}{\sqrt{1-k^2 \sin^2 t}}$ durch die Binomische Reihe darstelle.

22.6. Für die Bernoulli-Zahlen beweise man die asymptotische Beziehung

$$|B_{2k}| \sim 4\sqrt{\pi k} \left(\frac{k}{\pi e}\right)^{2k} \quad \text{für } k \rightarrow \infty.$$

22.7. Man zeige: Die Taylor-Entwicklung der Funktion $\frac{1}{\cos x}$ um den Nullpunkt hat die Gestalt

$$\frac{1}{\cos x} = \sum_{k=0}^{\infty} \frac{E_{2k}}{(2k)!} x^{2k}$$

mit positiven ganzen Zahlen E_{2k} (benannt nach Euler).

Man berechne $E_0, E_2, E_4, \dots, E_{10}$.

§ 23 Fourier-Reihen

In diesem letzten Paragraphen behandeln wir die wichtigsten Tatsachen aus der Theorie der Fourier-Reihen. Es handelt sich dabei um die Entwicklung von periodischen Funktionen nach dem Funktionensystem $\cos kx, \sin kx$, ($k \in \mathbb{N}$). Im Unterschied zu den Taylor-Reihen, die im Innern ihres Konvergenzbereichs immer gegen eine unendlich oft differenzierbare Funktion konvergieren, können durch Fourier-Reihen z.B. auch periodische Funktionen dargestellt werden, die nur stückweise stetig differenzierbar sind und deren Ableitungen Sprungstellen haben.

Periodische Funktionen

Eine auf ganz $\mathbb{R}$ definierte reell- oder komplexwertige Funktion f heißt *periodisch* mit der Periode $L > 0$, falls

$$f(x+L) = f(x) \quad \text{für alle } x \in \mathbb{R}.$$

Es gilt dann natürlich auch $f(x+nL) = f(x)$ für alle $x \in \mathbb{R}$ und $n \in \mathbb{Z}$. Durch eine Variablen-Transformation kann man Funktionen mit der Periode L auf solche mit der Periode 2π zurückführen: Hat f die Periode L , so hat die Funktion F , definiert durch

$$F(x) := f\left(\frac{L}{2\pi}x\right)$$

die Periode 2π . Aus der Funktion F kann man f durch die Formel

$$f(x) = F\left(\frac{2\pi}{L}x\right)$$

wieder zurückgewinnen. Bei der Behandlung periodischer Funktionen kann man sich also auf den Fall der Periode 2π beschränken. Im Folgenden verstehen wir unter periodischen Funktionen stets solche mit der Periode 2π .

Spezielle periodische Funktionen sind die *trigonometrischen Polynome*. Eine Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ heißt trigonometrisches Polynom der Ordnung n , falls sie sich schreiben läßt als

$$f(x) = \frac{a_0}{2} + \sum_{k=1}^n (a_k \cos kx + b_k \sin kx)$$

mit reellen Konstanten a_k, b_k . Die Konstanten sind durch die Funktion f ein-

deutig bestimmt, denn es gilt

$$a_k = \frac{1}{\pi} \int_0^{2\pi} f(x) \cos kx \, dx \quad \text{für } k = 0, 1, \dots, n,$$

$$b_k = \frac{1}{\pi} \int_0^{2\pi} f(x) \sin kx \, dx \quad \text{für } k = 1, \dots, n.$$

Dies folgt daraus, dass

$$\int_0^{2\pi} \cos kx \sin lx \, dx = 0 \quad \text{für alle natürlichen Zahlen } k \text{ und } l,$$

$$\int_0^{2\pi} \cos kx \cos lx \, dx = \int_0^{2\pi} \sin kx \sin lx \, dx = 0 \quad \text{für } k \neq l,$$

$$\int_0^{2\pi} \cos^2 kx \, dx = \int_0^{2\pi} \sin^2 kx \, dx = \pi \quad \text{für alle } k \geq 1.$$

Es ist häufig zweckmäßig, auch komplexwertige trigonometrische Polynome zu betrachten, bei denen für die Konstanten a_k, b_k beliebige komplexe Zahlen zugelassen sind. Unter Verwendung der Formeln

$$\cos x = \frac{1}{2}(e^{ix} + e^{-ix}), \quad \sin x = \frac{1}{2i}(e^{ix} - e^{-ix})$$

lässt sich das oben angegebene trigonometrische Polynom f auch schreiben als

$$f(x) = \sum_{k=-n}^n c_k e^{ikx},$$

wobei $c_0 = \frac{a_0}{2}$ und

$$c_k = \frac{1}{2}(a_k - ib_k), \quad c_{-k} = \frac{1}{2}(a_k + ib_k) \quad \text{für } k \geq 1.$$

Um in diesem Fall die Koeffizienten c_k durch Integration aus der Funktion f zu erhalten, brauchen wir den Begriff des Integrals einer komplexwertigen Funktion. Seien $u, v: [a, b] \rightarrow \mathbb{R}$ reelle Funktionen. Dann heißt die komplexwertige Funktion $\varphi := u + iv: [a, b] \rightarrow \mathbb{C}$ integrierbar, falls u und v integrierbar sind und

man setzt

$$\int_a^b (u(x) + iv(x)) dx := \int_a^b u(x) dx + i \int_a^b v(x) dx.$$

Speziell für die Funktion $\varphi(x) = e^{imx}$, $m \neq 0$, ergibt sich

$$\int_a^b e^{imx} dx = \frac{1}{im} e^{imx} \Big|_a^b,$$

also insbesondere

$$\int_0^{2\pi} e^{imx} dx = 0 \quad \text{für alle } m \in \mathbb{Z} \setminus \{0\}.$$

Damit erhält man für das trigonometrische Polynom $f(x) = \sum_{k=-n}^n c_k e^{ikx}$

$$c_k = \frac{1}{2\pi} \int_0^{2\pi} f(x) e^{-ikx} dx \quad \text{für } k = 0, \pm 1, \dots, \pm n,$$

da $f(x) e^{-ikx} = \sum_{m=-n}^n c_m e^{i(m-k)x}$.

Definition. Sei $f: \mathbb{R} \rightarrow \mathbb{C}$ eine periodische, über das Intervall $[0, 2\pi]$ integrierbare Funktion. Dann heißen die Zahlen

$$c_k := \frac{1}{2\pi} \int_0^{2\pi} f(x) e^{-ikx} dx, \quad k \in \mathbb{Z},$$

die *Fourier-Koeffizienten* von f , und die Reihe

$$\mathfrak{F}[f](x) := \sum_{k=-\infty}^{\infty} c_k e^{ikx},$$

d.h. die Folge der Partialsummen

$$\mathfrak{F}_n[f](x) := \sum_{k=-n}^n c_k e^{ikx}, \quad n \in \mathbb{N},$$

heißt *Fourier-Reihe* von f .

Die Fourier-Reihe lässt sich auch in der Form

$$\frac{a_0}{2} + \sum_{k=1}^{\infty} (a_k \cos kx + b_k \sin kx)$$

schreiben, wobei

$$a_k = \frac{1}{\pi} \int_0^{2\pi} f(x) \cos kx \, dx, \quad b_k = \frac{1}{\pi} \int_0^{2\pi} f(x) \sin kx \, dx.$$

Bemerkung. Ähnlich wie bei der Taylorreihe einer Funktion ist nicht garantiert, dass die Fourierreihe einer Funktion f konvergiert und dass sie im Falle der Konvergenz gegen f konvergiert.

Folgendes lässt sich aber leicht feststellen: Wenn die Funktion f sich überhaupt in der Gestalt

$$f(x) = \sum_{k=-\infty}^{\infty} \gamma_k e^{ikx}$$

mit *gleichmäßig* konvergenter Reihe darstellen lässt, dann muss diese Reihe die Fourier-Reihe von f sein. Weil nämlich gleichmäßige Konvergenz vorliegt, kann man bei der Berechnung der Fourier-Koeffizienten Integration und Limesbildung vertauschen und man erhält

$$\begin{aligned} c_k &= \frac{1}{2\pi} \int_0^{2\pi} \left(\sum_{m=-\infty}^{\infty} \gamma_m e^{imx} \right) e^{-ikx} \, dx \\ &= \frac{1}{2\pi} \sum_{m=-\infty}^{\infty} \int_0^{2\pi} \gamma_m e^{i(m-k)x} \, dx = \gamma_k. \end{aligned}$$

Im Allgemeinen konvergiert jedoch die Fourier-Reihe von f weder gleichmäßig noch punktweise gegen f . Den Fourier-Reihen ist ein anderer Konvergenzbegriff besser angepasst, die Konvergenz im quadratischen Mittel. Um diesen Begriff einzuführen, treffen wir zunächst einige Vorbereitungen.

Skalarprodukt für periodische Funktionen

Im Vektorraum V aller periodischen Funktionen $f: \mathbb{R} \rightarrow \mathbb{C}$, die über das Intervall $[0, 2\pi]$ Riemann-integrierbar sind, führen wir ein Skalarprodukt ein durch die Formel

$$\langle f, g \rangle := \frac{1}{2\pi} \int_0^{2\pi} \overline{f(x)} g(x) \, dx \quad \text{für } f, g \in V.$$

Folgende Eigenschaften sind leicht nachzuweisen ($f, g, h \in V, \lambda \in \mathbb{C}$):

- a) $\langle f+g, h \rangle = \langle f, h \rangle + \langle g, h \rangle,$
- b) $\langle f, g+h \rangle = \langle f, g \rangle + \langle f, h \rangle,$
- c) $\langle \lambda f, g \rangle = \bar{\lambda} \langle f, g \rangle,$
- d) $\langle f, \lambda g \rangle = \lambda \langle f, g \rangle,$
- e) $\langle f, g \rangle = \overline{\langle g, f \rangle}.$

Für jedes $f \in V$ gilt

$$\langle f, f \rangle = \frac{1}{2\pi} \int_0^{2\pi} |f(x)|^2 dx \geq 0.$$

Aus $\langle f, f \rangle = 0$ kann man jedoch i.Allg. nicht schließen, dass $f = 0$. Ist z.B. f im Intervall $[0, 2\pi]$ nur an endlich vielen Stellen von null verschieden, so gilt $\langle f, f \rangle = 0$. Für stetiges $f \in V$ folgt jedoch aus $\langle f, f \rangle = 0$, dass $f = 0$. Man setzt

$$\|f\|_2 := \sqrt{\langle f, f \rangle}.$$

Für diese Norm gilt die Dreiecksungleichung

$$\|f+g\|_2 \leq \|f\|_2 + \|g\|_2, \quad \text{vgl. (18.6).}$$

Definiert man die Funktion $e_k : \mathbb{R} \rightarrow \mathbb{C}$ durch

$$e_k(x) := e^{ikx},$$

so lassen sich die Fourier-Koeffizienten einer Funktion $f \in V$ einfach schreiben als

$$c_k = \langle e_k, f \rangle, \quad k \in \mathbb{Z}.$$

Die Funktionen e_k haben die Eigenschaft

$$\langle e_k, e_l \rangle = \delta_{kl} = \begin{cases} 0, & \text{falls } k \neq l, \\ 1, & \text{falls } k = l, \end{cases}$$

sie bilden also ein *Orthonormalsystem*.

Hilfssatz 1. Die Funktion $f \in V$ habe die Fourier-Koeffizienten $c_k, k \in \mathbb{Z}$. Dann gilt für alle $n \in \mathbb{N}$

$$\left\| f - \sum_{k=-n}^n c_k e_k \right\|_2^2 = \|f\|_2^2 - \sum_{k=-n}^n |c_k|^2.$$

Beweis. Wir setzen $g := \sum_{k=-n}^n c_k e_k$. Dann gilt

$$\langle f, g \rangle = \sum_{k=-n}^n c_k \langle f, e_k \rangle = \sum_{k=-n}^n c_k \bar{c}_k = \sum_{k=-n}^n |c_k|^2$$

und $\langle e_k, g \rangle = c_k$, also

$$\langle g, g \rangle = \sum_{k=-n}^n \bar{c}_k \langle e_k, g \rangle = \sum_{k=-n}^n |c_k|^2.$$

Daraus folgt

$$\begin{aligned} \|f - g\|_2^2 &= \langle f - g, f - g \rangle = \langle f, f \rangle - \langle f, g \rangle - \langle g, f \rangle + \langle g, g \rangle \\ &= \|f\|_2^2 - \sum_{k=-n}^n |c_k|^2 - \sum_{k=-n}^n |c_k|^2 + \sum_{k=-n}^n |c_k|^2 \\ &= \|f\|_2^2 - \sum_{k=-n}^n |c_k|^2, \quad \text{q.e.d.} \end{aligned}$$

Satz 1 (Besselsche Ungleichung). *Sei $f: \mathbb{R} \rightarrow \mathbb{C}$ eine periodische, über das Intervall $[0, 2\pi]$ Riemann-integrierbare Funktion mit den Fourier-Koeffizienten c_k . Dann gilt*

$$\sum_{k=-\infty}^{\infty} |c_k|^2 \leq \frac{1}{2\pi} \int_0^{2\pi} |f(x)|^2 dx.$$

Beweis. Aus Hilfssatz 1 folgt

$$\sum_{k=-n}^n |c_k|^2 \leq \|f\|_2^2$$

für alle $n \in \mathbb{N}$. Durch Grenzübergang ergibt sich die Behauptung.

Definition. Seien $f: \mathbb{R} \rightarrow \mathbb{C}$ und $f_n: \mathbb{R} \rightarrow \mathbb{C}$, $n \in \mathbb{N}$, periodische, über das Intervall $[0, 2\pi]$ Riemann-integrierbare Funktionen. Man sagt, die Folge (f_n) konvergiere im *quadratischen Mittel* gegen f , falls

$$\lim_{n \rightarrow \infty} \|f - f_n\|_2 = 0,$$

d.h. wenn das quadratische Mittel der Abweichung zwischen f und f_n , nämlich

$$\frac{1}{2\pi} \int_0^{2\pi} |f(x) - f_n(x)|^2 dx$$

für $n \rightarrow \infty$ gegen 0 konvergiert.

Man sieht unmittelbar: Konvergiert die Folge (f_n) gleichmäßig gegen f , so auch im quadratischen Mittel. Die Umkehrung gilt aber nicht. Eine im quadratischen Mittel konvergente Funktionenfolge braucht nicht einmal punktweise zu konvergieren.

Bemerkung. Der Hilfssatz 1 sagt, dass die Fourier-Reihe von f genau dann im quadratischen Mittel gegen f konvergiert, wenn

$$\sum_{k=-\infty}^{\infty} |c_k|^2 = \|f\|_2^2,$$

d.h. wenn die Besselsche Ungleichung zu einer Gleichung wird. Das Bestehen dieser Gleichung bezeichnet man auch als *Vollständigkeitsrelation*.

Hilfssatz 2. Sei $f: \mathbb{R} \rightarrow \mathbb{R}$ eine periodische Funktion, so dass $f|_{[0, 2\pi]}$ eine Treppenfunktion ist. Dann konvergiert die Fourier-Reihe von f im quadratischen Mittel gegen f .

Beweis

a) Wir behandeln zunächst den speziellen Fall, dass für f gilt

$$f(x) = \begin{cases} 1 & \text{für } 0 \leq x < a, \\ 0 & \text{für } a \leq x < 2\pi, \end{cases}$$

wobei a ein Punkt im Intervall $[0, 2\pi]$ ist. Die Fourier-Koeffizienten c_k dieser Funktion lauten

$$c_0 = \frac{a}{2\pi},$$

$$c_k = \frac{1}{2\pi} \int_0^a e^{-ikx} dx = \frac{i}{2\pi k} (e^{-ika} - 1) \quad \text{für } k \neq 0.$$

Für $k \neq 0$ gilt

$$|c_k|^2 = \frac{1}{4\pi^2 k^2} (1 - e^{ika}) (1 - e^{-ika}) = \frac{1 - \cos ka}{2\pi^2 k^2},$$

also

$$\begin{aligned} \sum_{k=-\infty}^{\infty} |c_k|^2 &= \frac{a^2}{4\pi^2} + \sum_{k=1}^{\infty} \frac{1 - \cos ka}{\pi^2 k^2} \\ &= \frac{a^2}{4\pi^2} + \frac{1}{\pi^2} \sum_{k=1}^{\infty} \frac{1}{k^2} - \frac{1}{\pi^2} \sum_{k=1}^{\infty} \frac{\cos ka}{k^2} \end{aligned}$$

$$= \frac{a^2}{4\pi^2} + \frac{1}{6} - \frac{1}{\pi^2} \left(\frac{(\pi-a)^2}{4} - \frac{\pi^2}{12} \right) = \frac{a}{2\pi},$$

wobei (21.9) benützt wurde. Es gilt deshalb

$$\sum_{k=-\infty}^{\infty} |c_k|^2 = \frac{a}{2\pi} = \frac{1}{2\pi} \int_0^{2\pi} |f(x)|^2 dx = \|f\|_2^2.$$

Nach Hilfssatz 1 folgt daraus die Konvergenz der Fourier-Reihe im quadratischen Mittel.

b) Ist $f| [0, 2\pi]$ eine beliebige Treppenfunktion, so gibt es Funktionen $f_1, \dots, f_r$ der in a) beschriebenen Gestalt und Konstanten $\gamma_1, \dots, \gamma_r$, so dass

$$f(x) = \sum_{j=1}^r \gamma_j f_j(x)$$

für alle $x \in \mathbb{R}$ mit evtl. Ausnahme der Sprungstellen. Für die n -ten Partialsummen $\mathfrak{F}_n[f]$ bzw. $\mathfrak{F}_n[f_j]$ der Fourierreihen von f und f_j gilt

$$\mathfrak{F}_n[f] = \sum_{j=1}^r \gamma_j \mathfrak{F}_n[f_j],$$

also

$$\|f - \mathfrak{F}_n[f]\|_2 = \left\| \sum_{j=1}^r \gamma_j (f_j - \mathfrak{F}_n[f_j]) \right\|_2 \leq \sum_{j=1}^r |\gamma_j| \cdot \|f_j - \mathfrak{F}_n[f_j]\|_2.$$

Nach Teil a) konvergiert dies für $n \rightarrow \infty$ gegen 0.

Satz 2. Sei $f: \mathbb{R} \rightarrow \mathbb{C}$ eine periodische Funktion, so dass $f| [0, 2\pi]$ Riemann-integrierbar ist. Dann konvergiert die Fourier-Reihe von f im quadratischen Mittel gegen f . Sind c_k die Fourier-Koeffizienten von f , so gilt die Vollständigkeitsrelation

$$\sum_{k=-\infty}^{\infty} |c_k|^2 = \frac{1}{2\pi} \int_0^{2\pi} |f(x)|^2 dx.$$

Beweis. Es genügt den Fall zu behandeln, dass f reellwertig ist und der Abschätzung $|f(x)| \leq 1$ für alle $x \in \mathbb{R}$ genügt.

Sei $\varepsilon > 0$ vorgegeben. Dann gibt es periodische Funktionen $\varphi, \psi: \mathbb{R} \rightarrow \mathbb{R}$ mit folgenden Eigenschaften:

a) $\varphi| [0, 2\pi]$ und $\psi| [0, 2\pi]$ sind Treppenfunktionen,

b) $-1 \leq \varphi \leq f \leq \psi \leq 1$,

c) $\int_0^{2\pi} (\psi(x) - \varphi(x)) dx \leq \frac{\pi}{4} \varepsilon^2$.

Wir setzen $g := f - \varphi$. Dann gilt

$$|g|^2 \leq |\psi - \varphi|^2 \leq 2(\psi - \varphi),$$

also

$$\frac{1}{2\pi} \int_0^{2\pi} |g(x)|^2 dx \leq \frac{1}{\pi} \int_0^{2\pi} (\psi(x) - \varphi(x)) dx \leq \frac{\varepsilon^2}{4}.$$

Für die Partialsummen $\mathfrak{F}_n[f]$, $\mathfrak{F}_n[\varphi]$ bzw. $\mathfrak{F}_n[g]$ der Fourier-Reihen von f , φ bzw. g gilt

$$\mathfrak{F}_n[f] = \mathfrak{F}_n[\varphi] + \mathfrak{F}_n[g].$$

Nach Hilfssatz 2 gibt es ein N , so dass

$$\|\varphi - \mathfrak{F}_n[\varphi]\|_2 \leq \frac{\varepsilon}{2} \quad \text{für alle } n \geq N.$$

Für alle n gilt nach Hilfssatz 1

$$\|g - \mathfrak{F}_n[g]\|_2^2 \leq \|g\|_2^2 \leq \frac{\varepsilon^2}{4}.$$

Daher gilt für alle $n \geq N$

$$\|f - \mathfrak{F}_n[f]\|_2 \leq \|\varphi - \mathfrak{F}_n[\varphi]\|_2 + \|g - \mathfrak{F}_n[g]\|_2 \leq \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon,$$

die Fourier-Reihe konvergiert also im quadratischen Mittel gegen f . Wie schon bemerkt, folgt daraus, dass aus der Besselschen Ungleichung eine Gleichung wird.

Bemerkung. Schreibt man die Fourier-Reihe einer periodischen Funktion f in der Form

$$\frac{a_0}{2} + \sum_{k=1}^{\infty} (a_k \cos kx + b_k \sin kx),$$

so lautet die Vollständigkeitsrelation

$$\frac{1}{2}|a_0|^2 + \sum_{k=1}^{\infty} (|a_k|^2 + |b_k|^2) = \frac{1}{\pi} \int_0^{2\pi} |f(x)|^2 dx.$$

Dies ergibt sich durch einfaches Umrechnen der Koeffizienten c_k in a_k und b_k .

Beispiele

(23.1) Wir betrachten die schon in Beispiel (21.2) untersuchte periodische Funktion $\sigma: \mathbb{R} \rightarrow \mathbb{R}$ mit

$$\sigma(x) = \frac{\pi - x}{2} \quad \text{für } 0 < x < 2\pi, \quad \sigma(0) = 0.$$

Die Berechnung der Fourier-Koeffizienten ergibt

$$c_0 = \frac{1}{2\pi} \int_0^{2\pi} \sigma(x) dx = 0$$

und für $k \neq 0$

$$\begin{aligned} c_k &= \frac{1}{2\pi} \int_0^{2\pi} \frac{\pi - x}{2} e^{-ikx} dx = -\frac{1}{4\pi} \int_0^{2\pi} x e^{-ikx} dx = [\text{Partielle Integration}] \\ &= \frac{1}{4ik\pi} \left(x e^{-ikx} \Big|_0^{2\pi} - \underbrace{\int_0^{2\pi} e^{-ikx} dx}_{=0} \right) = \frac{1}{2ik}. \end{aligned}$$

Die Fourier-Reihe von f lautet daher

$$\sum_{k=1}^{\infty} \frac{1}{2i} \left(\frac{e^{ikx}}{k} - \frac{e^{-ikx}}{k} \right) = \sum_{k=1}^{\infty} \frac{\sin kx}{k}.$$

Damit haben wir die Reihe wiedergefunden, von der wir bereits in (21.2) gezeigt haben, dass sie überall punktweise und in jedem Intervall $[\delta, 2\pi - \delta]$, ($0 < \delta < \pi$), gleichmäßig gegen $\sigma(x)$ konvergiert. Einige Partialsummen der Fourier-Reihe sind in Bild 23.1 dargestellt.

Nach Satz 2 konvergiert die Reihe auch im quadratischen Mittel gegen σ und die Vollständigkeitsrelation liefert

$$\sum_{k=1}^{\infty} \frac{1}{k^2} = \frac{1}{\pi} \int_0^{2\pi} |\sigma(x)|^2 dx = \frac{1}{\pi} \int_0^{2\pi} \left| \frac{\pi - x}{2} \right|^2 dx = \frac{1}{4\pi} \int_{-\pi}^{\pi} x^2 dx = \frac{\pi^2}{6},$$

was uns schon aus (21.9) bekannt ist.

(23.2) Wir haben in (21.9) hergeleitet, dass für $0 \leq x \leq 2\pi$ gilt

$$\frac{(\pi - x)^2}{4} = \frac{\pi^2}{12} + \sum_{k=1}^{\infty} \frac{\cos kx}{k^2} = \frac{\pi^2}{12} + \sum_{k=1}^{\infty} \frac{1}{2k^2} (e^{ikx} + e^{-ikx}). \quad (*)$$

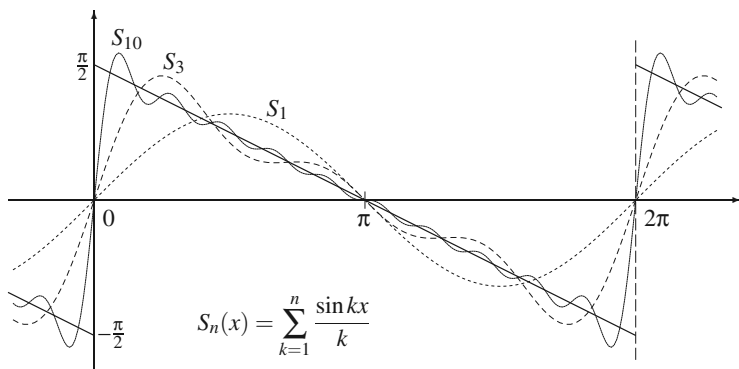


Bild 23.1 Zur Fourier-Entwicklung der Funktion σ

Die Reihe (*) konvergiert gleichmäßig, stellt also die Fourier-Reihe derjenigen periodischen Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ dar, die für $x \in [0, 2\pi]$ mit $(\pi - x)^2/4$ übereinstimmt, vgl. Bild 23.2.

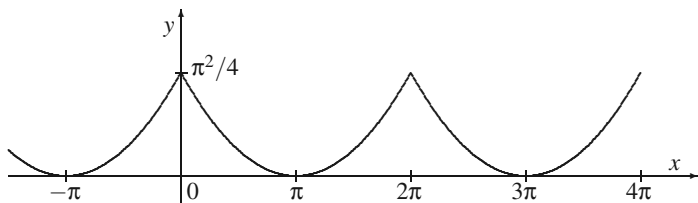


Bild 23.2 Die periodisch fortgesetzte Funktion $y = \frac{(\pi-x)^2}{4}$, $0 \leq x \leq 2\pi$

Die Vollständigkeitsrelation liefert

$$\frac{\pi^4}{144} + 2 \sum_{k=1}^{\infty} \frac{1}{4k^4} = \frac{1}{2\pi} \int_0^{2\pi} \frac{(\pi-x)^4}{16} dx = \frac{\pi^4}{80},$$

also

$$\sum_{k=1}^{\infty} \frac{1}{k^4} = \frac{\pi^4}{90}$$

in Übereinstimmung mit Satz 11 aus §22.

Die Formeln für die Fourier-Reihen $\sum_1^\infty \frac{\sin kx}{k}$ und $\sum_1^\infty \frac{\cos kx}{k^2}$ aus (23.1) und (23.2) sind nur die ersten Glieder einer ganzen Folge von Fourier-Reihen, die mit den sog. Bernoulli-Polynomen zusammenhängen. Diese besprechen wir jetzt.

(23.3) Bernoulli-Polynome. Für $n \in \mathbb{N}$ ist das n -te Bernoulli-Polynom ist definiert als

$$B_n(x) := \sum_{k=0}^n \binom{n}{k} B_k x^{n-k},$$

wobei B_k die in (22.10) definierten Bernoulli-Zahlen sind. $B_n(x)$ ist ein Polynom n -ten Grades mit $B_n(0) = B_n$. Die Bernoulli-Polynome vom Grad ≤ 6 lauten:

$$B_0(x) = 1,$$

$$B_1(x) = x - \frac{1}{2},$$

$$B_2(x) = x^2 - x + \frac{1}{6},$$

$$B_3(x) = x^3 - \frac{3}{2}x^2 + \frac{1}{2}x,$$

$$B_4(x) = x^4 - 2x^3 + x^2 - \frac{1}{30},$$

$$B_5(x) = x^5 - \frac{5}{2}x^4 + \frac{5}{3}x^3 - \frac{1}{6}x,$$

$$B_6(x) = x^6 - 3x^5 + \frac{5}{2}x^4 - \frac{1}{2}x^2 + \frac{1}{42}.$$

Wir zeigen jetzt folgende Eigenschaften der Bernoulli-Polynome:

- i) $B'_n(x) = nB_{n-1}(x)$ für alle $n \geq 1$,
- ii) $B_n(0) = B_n(1)$ für alle $n \geq 2$,
- iii) $\int_0^1 B_n(x) dx = 0$ für alle $n \geq 1$.

Beweis. Zu i) Unter Benutzung der Gleichung $(n-k)\binom{n}{k} = n\binom{n-1}{k}$ erhält man

$$\begin{aligned} B'_n(x) &= \sum_{k=0}^{n-1} (n-k) \binom{n}{k} B_k x^{n-1-k} \\ &= n \sum_{k=0}^{n-1} \binom{n-1}{k} B_k x^{n-1-k} = nB_{n-1}(x). \end{aligned}$$

Zu ii) Die Rekursionsformel für die Bernoulli-Zahlen (22.10) ist äquivalent zu $\sum_{k=0}^{n-1} \binom{n}{k} B_k = 0$ für $n \geq 2$. Daraus folgt

$$B_n(1) = \sum_{k=0}^{n-1} \binom{n}{k} B_k + B_n = B_n = B_n(0) \quad \text{für } n \geq 2.$$

Zu iii)

$$(n+1) \int_0^1 B_n(x) dx = \int_0^1 B'_{n+1}(x) dx = (B_{n+1}(1) - B_{n+1}(0) = 0, \text{ q.e.d.}$$

Wir definieren jetzt für $n \geq 1$ periodische Funktionen $\tilde{B}_n: \mathbb{R} \rightarrow \mathbb{R}$ mit der Periode 1 durch

$$\begin{aligned} \tilde{B}_1(x) &= B_1(x) \quad \text{für } 0 < x < 1, \quad \tilde{B}_1(0) = 0, \\ \tilde{B}_n(x) &= B_n(x) \quad \text{für } 0 \leq x < 1, \quad (n \geq 2) \end{aligned}$$

und

$$\tilde{B}_n(x+m) = \tilde{B}_n(x) \quad \text{für alle } x \in \mathbb{R}, m \in \mathbb{Z}, n \geq 1.$$

Die Funktionen $\tilde{B}_n$ sind für $n \geq 2$ stetig und $(n-2)$ -mal stetig differenzierbar.

Durch die Variablen-Transformation $x \mapsto 2\pi x$ erhält man aus den Beispielen (23.1) und (23.2) die Formeln

$$\tilde{B}_1(x) = -2 \sum_{m=1}^{\infty} \frac{\sin(2\pi m x)}{2\pi m} \quad \text{und} \quad \tilde{B}_2(x) = 4 \sum_{m=1}^{\infty} \frac{\cos(2\pi m x)}{(2\pi m)^2}.$$

Dies verallgemeinert sich wie folgt:

$$\begin{aligned} \tilde{B}_{2k+1}(x) &= (-1)^{k-1} 2(2k+1)! \sum_{m=1}^{\infty} \frac{\sin(2\pi m x)}{(2\pi m)^{2k+1}}, \quad (k \geq 0), \\ \tilde{B}_{2k}(x) &= (-1)^{k-1} 2(2k)! \sum_{m=1}^{\infty} \frac{\cos(2\pi m x)}{(2\pi m)^{2k}}, \quad (k \geq 1). \end{aligned}$$

Beweis. Wir bezeichnen für $n = 2k+1$ bzw. $n = 2k$ die rechten Seiten der Formeln mit $\beta_n(x)$. Man sieht unmittelbar

$$\beta'_{n+1}(x) = (n+1)\beta_n(x) \quad \text{und} \quad \int_0^1 \beta_n(x) dx = 0$$

für alle $n \geq 2$ (da gliedweise Differentiation bzw. Integration erlaubt ist). Dieselben Beziehungen gelten für die Funktionen $\tilde{B}_n(x)$. Da $\tilde{B}_n(x) = \beta_n(x)$ für $n = 1, 2$, folgt durch vollständige Induktion über n , dass $\tilde{B}_n(x) = \beta_n(x)$ für alle $n \geq 1$, q.e.d.

Bemerkung. Setzt man in der Formel für $\tilde{B}_{2k}(x)$ die Variable $x = 0$, so erhält man wieder die Werte für $\zeta(2k)$ aus §22, Satz 11.

Setzt man dagegen $x = \frac{1}{4}$ in die Formel für $\widetilde{B}_{2k+1}(x)$, so ergibt sich als Verallgemeinerung der Leibniz'schen Reihe (Fall $k = 0$)

$$\sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)^{2k+1}} = (-1)^{k-1} \frac{(2\pi)^{2k+1}}{2(2k+1)!} B_{2k+1}\left(\frac{1}{4}\right).$$

Für $k = 1, 2$ erhält man wegen $B_3(\frac{1}{4}) = \frac{3}{64}$ und $B_5(\frac{1}{4}) = -\frac{25}{1024}$

$$\sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)^3} = \frac{\pi^3}{32} \quad \text{und} \quad \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)^5} = \frac{5\pi^5}{1536}.$$

Wir kommen nun zu einer großen Klasse von Funktionen, deren Fourier-Reihe gleichmäßig konvergiert.

Satz 3. *Es sei $f: \mathbb{R} \rightarrow \mathbb{C}$ eine stetige periodische Funktion, die stückweise stetig differenzierbar ist, d.h. es gebe eine Unterteilung*

$$0 = t_0 < t_1 < \dots < t_r = 2\pi$$

von $[0, 2\pi]$, so dass $f|_{[t_{k-1}, t_k]}$ für $k = 1, \dots, r$ stetig differenzierbar ist. Dann konvergiert die Fourier-Reihe von f gleichmäßig gegen f .

Ein Beispiel für Satz 3 ist die in (23.2) untersuchte Funktion.

Beweis. Es sei $\varphi_k: [t_{k-1}, t_k] \rightarrow \mathbb{C}$ die stetige Ableitung von $f|_{[t_{k-1}, t_k]}$ und $\varphi: \mathbb{R} \rightarrow \mathbb{C}$ diejenige periodische Funktion, die auf $[t_{k-1}, t_k]$ mit φ_k übereinstimmt. Für die Fourier-Koeffizienten γ_n der Funktion φ gilt nach der Besselschen Ungleichung

$$\sum_{n=-\infty}^{\infty} |\gamma_n|^2 \leq \|\varphi\|_2^2 < \infty.$$

Für $n \neq 0$ lassen sich die Fourier-Koeffizienten c_n von f wie folgt durch partielle Integration aus den Fourier-Koeffizienten γ_n von φ gewinnen:

$$\begin{aligned} \int_{t_{k-1}}^{t_k} f(x) e^{-inx} dx &= \frac{i}{n} \int_{t_{k-1}}^{t_k} f(x) d(e^{-inx}) \\ &= \frac{i}{n} \left(f(x) e^{-inx} \Big|_{t_{k-1}}^{t_k} - \int_{t_{k-1}}^{t_k} \varphi(x) e^{-inx} dx \right). \end{aligned}$$

Da wegen der Periodizität von f

$$\sum_{k=1}^r \left(f(x) e^{-inx} \Big|_{t_{k-1}}^{t_k} \right) = -f(t_0) e^{-int_0} + f(t_r) e^{-int_r} = 0,$$

folgt

$$\begin{aligned} c_n &= \frac{1}{2\pi} \int_0^{2\pi} f(x) e^{-inx} dx = \frac{1}{2\pi} \sum_{k=1}^r \int_{t_{k-1}}^{t_k} f(x) e^{-inx} dx \\ &= \frac{-i}{2\pi n} \int_0^{2\pi} \varphi(x) e^{-inx} dx = \frac{-i\gamma_n}{n}. \end{aligned}$$

Wegen der für alle $a, b \in \mathbb{R}$ gültigen Ungleichung $ab \leq \frac{1}{2}(a^2 + b^2)$ ergibt sich

$$|c_n| = \frac{|\gamma_n|}{|n|} \leq \frac{1}{2} \left(\frac{1}{|n|^2} + |\gamma_n|^2 \right).$$

Weil $\sum_{n=1}^{\infty} \frac{1}{n^2}$ und $\sum_{n=-\infty}^{\infty} |\gamma_n|^2$ konvergent sind, folgt

$$\sum_{n=-\infty}^{\infty} |c_n| < \infty.$$

Die Fourier-Reihe $\sum_{n=-\infty}^{\infty} c_n e^{inx}$ von f konvergiert also absolut und gleichmäßig gegen eine (nach §21, Satz 1) stetige Funktion g . Somit konvergiert die Fourier-Reihe im quadratischen Mittel sowohl gegen f als auch gegen g , woraus folgt

$$\|f - g\|_2 = 0.$$

Da f und g stetig sind, folgt daraus, dass f und g übereinstimmen. Satz 3 ist damit bewiesen.

Bemerkung. Satz 3 lässt sich auf unstetige Funktionen verallgemeinern, vgl. Aufgabe 23.6.

(23.4) Die Fresnelschen Integrale. Als eine Anwendung von Satz 3 berechnen wir die uneigentlichen Integrale

$$\int_0^{\infty} \cos(x^2) dx = \int_0^{\infty} \sin(x^2) dx = \frac{\sqrt{\pi}}{2\sqrt{2}}.$$

Diese Integrale spielen in der Optik in der Fresnelschen Beugungstheorie eine Rolle.

Man kann sich leicht überlegen, dass die Integrale konvergieren, vgl. Aufgabe 20.6. Zu ihrer Berechnung betrachten wir folgende Funktion $F : [0, 2\pi] \rightarrow \mathbb{C}$

$$F(x) := e^{ix^2/2\pi} = \cos\left(\frac{x^2}{2\pi}\right) + i \sin\left(\frac{x^2}{2\pi}\right).$$

Da $F(0) = F(2\pi) = 1$, lässt sich F zu einer auf ganz $\mathbb{R}$ stetigen und stückweise stetig differenzierbaren periodischen Funktion $F : \mathbb{R} \rightarrow \mathbb{C}$ fortsetzen, deren Fourier-Reihe

$$F(x) = \sum_{n \in \mathbb{Z}} c_n e^{inx}$$

nach Satz 3 gleichmäßig gegen F konvergiert. Insbesondere gilt $\sum_{n \in \mathbb{Z}} c_n = 1$.

Die Fourier-Koeffizienten von F sind

$$c_n = \frac{1}{2\pi} \int_0^{2\pi} e^{ix^2/2\pi} e^{-inx} dx.$$

Nun ist

$$\begin{aligned} e^{ix^2/2\pi} e^{-inx} &= \exp\left(\frac{i}{2\pi}(x^2 - 2\pi nx)\right) \\ &= \exp\left(\frac{i}{2\pi}(x - \pi n)^2\right) \exp\left(-\frac{i\pi}{2}n^2\right) \\ &= \mathfrak{v}_n \exp\left(\frac{i}{2\pi}(x - \pi n)^2\right) \end{aligned}$$

mit

$$\mathfrak{v}_n := \begin{cases} 1, & \text{falls } n \text{ gerade,} \\ -i, & \text{falls } n \text{ ungerade.} \end{cases}$$

Es folgt

$$c_n = \frac{\mathfrak{v}_n}{2\pi} \int_0^{2\pi} e^{i(x-\pi n)^2/2\pi} dx = \frac{\mathfrak{v}_n}{2\pi} \int_{-\pi n}^{(2-n)\pi} e^{ix^2/2\pi} dx.$$

Summation über alle geraden und alle ungeraden n ergibt

$$1 = \sum_{n \in \mathbb{Z}} c_n = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{ix^2/2\pi} dx - \frac{i}{2\pi} \int_{-\infty}^{\infty} e^{ix^2/2\pi} dx,$$

also

$$\int_{-\infty}^{\infty} e^{ix^2/2\pi} dx = \frac{2\pi}{1-i} = \pi(1+i).$$

Die Variablen-Substitution $t = x/\sqrt{2\pi}$ liefert

$$\int_{-\infty}^{\infty} e^{it^2} dt = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{ix^2/2\pi} dx = \sqrt{\pi} \frac{1+i}{\sqrt{2}},$$

woraus sich durch Trennung in Real- und Imaginärteil die behaupteten Werte der Fresnelschen Integrale ergeben.

AUFGABEN

23.1. Man berechne die Fourier-Reihe der periodischen Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ mit

$$f(x) = |x| \quad \text{für } -\pi \leq x \leq \pi.$$

23.2. Man berechne die Fourier-Reihe der Funktion $f(x) = |\sin x|$.

23.3. Man beweise: Ist $f: \mathbb{R} \rightarrow \mathbb{R}$ eine gerade (bzw. ungerade) periodische Funktion, so hat die Fourier-Reihe von f die Gestalt

$$\frac{a_0}{2} + \sum_{k=1}^{\infty} a_k \cos kx \quad \text{bzw.} \quad \sum_{k=1}^{\infty} b_k \sin kx.$$

23.4. a) Man zeige: Jede stetige periodische Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ lässt sich gleichmäßig durch stetige, stückweise lineare periodische Funktionen approximieren. Dabei heißt eine stetige periodische Funktion $\varphi: \mathbb{R} \rightarrow \mathbb{R}$ stückweise linear, wenn die Funktion $f|_{[0, 2\pi]}$ stückweise linear im Sinne der Definition in Aufgabe 11.8 ist.

b) Man beweise mit Teil a) und Satz 3, dass sich jede stetige periodische Funktion $f: \mathbb{R} \rightarrow \mathbb{C}$ gleichmäßig durch trigonometrische Polynome approximieren lässt (Weierstraß'scher Approximationssatz für periodische Funktionen).

23.5. Sei $f: \mathbb{R} \rightarrow \mathbb{C}$ eine stetige periodische Funktion mit Fourier-Koeffizienten c_n , $n \in \mathbb{Z}$. Man beweise:

a) Ist f k -mal stetig differenzierbar, so folgt

$$c_n = O\left(\frac{1}{|n|^k}\right) \quad \text{für } |n| \rightarrow \infty.$$

b) Falls

$$c_n = O\left(\frac{1}{|n|^{k+2}}\right) \quad \text{für } |n| \rightarrow \infty,$$

so ist f k -mal stetig differenzierbar und die Fourier-Reihe konvergiert gleichmäßig gegen f .

23.6. Die periodische (nicht notwendig stetige) Funktion $f: \mathbb{R} \rightarrow \mathbb{C}$ sei stückweise stetig differenzierbar, d.h. es gebe eine Unterteilung

$$0 = t_0 < t_1 < \dots < t_{r-1} < t_r = 2\pi,$$

so dass sich die Funktionen $f|_{]t_{j-1}, t_j[}$ zu stetig differenzierbaren Funktionen $f_j : [t_{j-1}, t_j] \rightarrow \mathbb{C}$ fortsetzen lassen ($j = 1, \dots, r$). Es seien

$$f_+(t_j) := \lim_{t \searrow t_j} f(t) \quad \text{und} \quad f_-(t_j) := \lim_{t \nearrow t_j} f(t)$$

die rechts- bzw. linksseitigen Grenzwerte von f an den Stellen t_j und

$$\gamma_j := f_+(t_j) - f_-(t_j)$$

die Sprunghöhen von f an diesen Stellen. Man beweise:

a) Die Funktion $F : \mathbb{R} \rightarrow \mathbb{C}$,

$$F(t) := f(t) - \sum_{j=1}^r \frac{\gamma_j}{\pi} \sigma(t - t_j),$$

wobei $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ die in Beispiel (23.1) betrachtete Funktion ist, stetig und stückweise stetig differenzierbar.

b) Die Fourier-Reihe von f konvergiert auf jedem kompakten Intervall $[a, b] \subset \mathbb{R}$, das keine Unstetigkeitsstelle von f enthält, gleichmäßig gegen f . An den Stellen t_j konvergiert die Fourier-Reihe von f gegen den Mittelwert

$$\frac{1}{2}(f_+(t_j) + f_-(t_j)).$$

23.7. Sei $a \in \mathbb{R} \setminus \mathbb{Z}$ und $f : \mathbb{R} \rightarrow \mathbb{C}$ die periodische Funktion mit

$$f(x) = e^{iax} \quad \text{für } 0 \leq x < 2\pi, \quad f(x + 2\pi n) = f(x), \quad (n \in \mathbb{Z}).$$

Man berechne die Fourier-Reihe von f und bestimme ihr Konvergenzverhalten (vgl. Aufgabe 23.6).

Was ergibt sich für $x = 0$?

23.8. In dieser Aufgabe werden die Bernoulli-Polynome aus (23.3) benutzt.

Man zeige:

a) Sei $f : [0, 1] \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion. Dann gilt

$$\frac{1}{2}(f(0) + f(1)) = \int_0^1 f(x) dx + \int_0^1 B_1(x) f'(x) dx.$$

b) Ist $f : [0, 1] \rightarrow \mathbb{R}$ sogar $2r$ -mal stetig differenzierbar ($r \geq 1$), so gilt

$$\begin{aligned} \int_0^1 B_1(x) f'(x) dx &= \sum_{j=1}^r \frac{B_{2j}}{(2j)!} \left(f^{(2j-1)}(1) - f^{(2j-1)}(0) \right) \\ &\quad - \int_0^1 \frac{B_{2r}(x)}{(2r)!} f^{(2r)}(x) dx. \end{aligned}$$

c) Seien $m < n$ ganze Zahlen und $f: [m, n] \rightarrow \mathbb{R}$ eine $2r$ -mal stetig differenzierbare Funktion. Dann gilt die *Euler-MacLaurinsche* Summationsformel

$$\sum_{k=m}^n f(k) = \frac{1}{2}(f(m) + f(n)) + \int_m^n f(x) dx + \sum_{j=1}^r \frac{B_{2j}}{(2j)!} (f^{(2j-1)}(n) - f^{(2j-1)}(m)) + R_{2r}$$

mit

$$R_{2r} = - \int_m^n \frac{\tilde{B}_{2r}(x)}{(2r)!} f^{(2r)}(x) dx \quad \text{und} \quad |R_{2r}| \leq \frac{|B_{2r}|}{(2r)!} \int_m^n |f^{(2r)}(x)| dx.$$

23.9. Zur Berechnung der Euler-Mascheronischen Konstanten

$$\gamma = \lim_{N \rightarrow \infty} \left(\sum_{n=1}^N \frac{1}{n} - \log N \right)$$

werte man für $M \geq 1$ und $r \geq 1$ den Limes

$$\lim_{N \rightarrow \infty} \left(\sum_{n=M}^N \frac{1}{n} - \int_M^N \frac{dx}{x} \right)$$

mithilfe der Euler-MacLaurinschen Summationsformel aus und beweise die Näherungs-Formel

$$\gamma = \left(\sum_{k=1}^M \frac{1}{k} - \log M \right) - \frac{1}{2M} + \sum_{j=1}^{r-1} \frac{B_{2j}}{2j} \cdot \frac{1}{M^{2j}} + \theta \cdot \left(\frac{B_{2r}}{2r} \cdot \frac{1}{M^{2r}} \right)$$

mit $0 \leq \theta \leq 1$.

i) Durch geeignete Wahl von M und r berechne man γ auf 50 Dezimalstellen genau.

ii) Für jedes feste $M \geq 1$ gilt $\lim_{r \rightarrow \infty} \left(\frac{B_{2r}}{2r} \cdot \frac{1}{M^{2r}} \right) = 0$.

iii) Wie kann man M und r wählen, um γ auf 1000 Dezimalstellen genau zu berechnen?

Zusammenstellung der Axiome der reellen Zahlen

Körperaxiome: Es sind zwei Verknüpfungen (Addition und Multiplikation) definiert, so dass folgende Axiome erfüllt sind:		Körper	angeordneter Körper	archimedisch angeordneter Körper	vollständiger archimedisch angeordneter Körper
Axiome der Addition Assoziativgesetz Kommutativgesetz Existenz der Null Existenz des Negativen	Axiome der Multiplikation Assoziativgesetz Kommutativgesetz Existenz der Eins ($\neq 0$) Existenz des Inversen (zu Elementen $\neq 0$)				
Distributivgesetz					
Anordnungsaxiome: Es sind gewisse Elemente als positiv ausgezeichnet ($x > 0$), so dass folgende Axiome erfüllt sind:					
Für jedes Element x gilt genau eine der Beziehungen $x > 0, x = 0, -x > 0$					
$x > 0 \wedge y > 0 \Rightarrow x + y > 0$					
$x > 0 \wedge y > 0 \Rightarrow xy > 0$					
Archimedisches Axiom: Zu $x > 0, y > 0$ existiert eine natürliche Zahl n mit $nx > y$.					
Vollständigkeitsaxiom: Jede Cauchy-Folge konvergiert.					

Literaturhinweise

Einführungen in die Analysis

- [AE] H. Amann und J. Escher: Analysis I. Birkhäuser, 3. Aufl. 2006.
- [BF] M. Barner und F. Flohr: Analysis I. De Gruyter, 5. Aufl. 2000.
- [Brö] Th. Bröcker: Analysis 1. Spektrum, Akad. Verl., 2. Aufl. 1995.
- [FW] O. Forster und R. Wessoly: Übungsbuch zur Analysis 1. Springer Spektrum, 6. Aufl. 2012.
- [He] H. Heuser: Lehrbuch der Analysis, Teil 1. Vieweg+Teubner, 17. Aufl. 2009.
- [Ho] H.S. Holdgrün: Analysis, Band 1. Leins Verlag Göttingen 1998.
- [HW] E. Hairer und G. Wanner: Analysis in historischer Entwicklung. Springer 2011.
- [Ka] W. Kabblo: Einführung in die Analysis I. Spektrum, Akad. Verl., 2. Aufl. 2000.
- [Kö] K. Königsberger: Analysis 1. Springer, 6. Aufl. 2004.
- [Wa] W. Walter, Analysis I. Springer, 7. Aufl. 2004.

Weitere im Text zitierte Literatur

- [AH] J. Arndt und Ch. Haenel: Pi. Algorithmen, Computer, Arithmetik. Springer, 2. Aufl. 2000.
- [BM] R. Braun und R. Meise: Analysis mit Maple. Vieweg+Teubner, 2. Aufl. 2011.
- [Br] D.S. Bridges: Computability. A Mathematical Sketchbook. Springer 2007.
- [FB] E. Freitag und R. Busam: Funktionentheorie 1. Springer, 4. Aufl. 2007.
- [Fi] G. Fischer: Lineare Algebra. Vieweg+Teubner, 17. Aufl. 2009.
- [FL] W. Fischer und I. Lieb: Funktionentheorie. Vieweg, 9. Aufl. 2005.
- [Fo] O. Forster: Algorithmische Zahlentheorie. Vieweg 1996.
- [H] D. Hilbert: Grundlagen der Geometrie, Teubner, 14. Aufl. 1999.
- [HU] J.E. Hopcroft und J.D. Ullman: Einführung in die Automatentheorie, Formale Sprachen und Komplexitätstheorie. Oldenbourg, 2. Aufl. 2002.

- [HWr] G.H. Hardy and E.M. Wright: An introduction to the theory of numbers. Oxford U.P., 6th ed. 2008.
- [L] E. Landau: Grundlagen der Analysis. Teubner 1930. Nachdruck Heldermann 2004.
- [Ri] P. Ribenboim: Meine Zahlen, meine Freunde. Springer 2009.
- [SG] H. Stoppel und B. Gries: Übungsbuch zur Linearen Algebra. Vieweg+Teubner, 7. Aufl. 2010.
- [T] F. Toenniessen: Das Geheimnis der transzendenten Zahlen. Spektrum, Akad. Verl. 2010.
- [Z] H.D. Ebbinghaus et al.: Zahlen. Springer, 3. Aufl. 2010.

Namens- und Sachverzeichnis

- Abel, Niels Henrik* (1802–1829), 286
 Abelscher Grenzwertsatz, 286
 abgeschlossenes Intervall, 90
 Ableitung, 163
 Ableitung höherer Ordnung, 174
 absolut konvergent, 74
 Absolut-Betrag, 29
 abzählbar, 91
 Additionstheorem der Exponentialfunktion, 87
 Additionstheorem des Tangens, 160
 Additionstheoreme von Sinus und Cosinus, 146
 allgemeine Potenz, 128
 alternierende harmonische Reihe, 74, 216, 242
 alternierende Reihen, 72
 angeordneter Körper, 28
 Anordnungs-Axiome, 25
Archimedes (287?–212 v.Chr.), 31
 Archimedisches Axiom, 31
 Arcus-Cosinus, 154
 Arcus-Sinus, 155
 Arcus-Tangens, 155
 Arcus-Tangens-Reihe, 288
 Area cosinus hyperbolici, 134
 Area sinus hyperbolici, 134
 Argument einer komplexen Zahl, 157
 arithmetisch-geometrisches Mittel, 67
 arithmetisches Mittel, 67
 Assoziativgesetz, 17, 19
 asymptotisch gleich, 248

b-adischer Bruch, 52
 berechenbare Zahl, 93
Bernoulli, Jacob (1654–1705), 32
 Bernoulli-Polynome, 315
 Bernoulli-Zahlen, 298
 Bernoullische Ungleichung, 32
 Berührungspunkt, 94
 beschränkt, 38
 beschränkte Folgen, 38
 beschränkte Funktion, 117
 beschränkte Menge, 95
Bessel, Friedrich Wilhelm (1784–1846), 309
 Besselsche Ungleichung, 309
 bestimmte Divergenz, 45
 Beta-Funktion, 254
 Betrag, 29
 Betrag einer komplexen Zahl, 138
 bewerteter Körper, 30
 Binär-Darstellung, 58
Binomi, Alessandro (1727–1643), 9
 Binomial-Koeffizienten, 7
 Binomische Reihe, 289
 Binomischer Lehrsatz, 9
Bohr, Harald (1887–1951), 244
 Satz von Bohr/Møllerup, 244
Bolzano, Bernhard (1781–1848), 55
 Satz von Bolzano-Weierstraß, 55

Cantor, Georg (1845–1918), 93
 Cantorsches Diagonalverfahren, 93
Cauchy, Augustin Louis (1789–1857), 49
 Cauchy-Folge, 49
 Cauchy-Produkt von Reihen, 85
 Cauchy-Schwarz'sche Ungleichung, 185
 Cauchysches Konvergenz-Kriterium, 70
 ceil, 32
Cohen, Paul J. (1934–2007), 101

- Cosecans, 154
Cosinus, 145
Cosinus hyperbolicus, 111
Cotangens, 153
 Partialbruch-Zerlegung, 271
de l'Hospital, siehe Hospital, 187
Dedekind, Richard (1835–1916), 102
Dedekindsches Schnittaxiom, 102
Definition durch vollständige Induktion, 2
Dezimalbruch, 53
 periodischer, 45
Differentialquotient, 163
differenzierbar, 163
 von links, 167
 von rechts, 167
Dirichlet, Peter Gustav Lejeune (1805–1859), 105
Dirichletsche Funktion, 105, 205
Distributivgesetz, 19
divergent, 37
divergent, bestimmt, 45
Doppelsummen, 21
Dreiecks-Ungleichung, 30, 138

Einheitswurzeln, 158
Einselement, 19
Einstein, Albert (1879–1955), 292
Einsteinsche Gleichung $E = mc^2$, 292
 ε -Umgebung, 36
 ε - δ -Definition der Stetigkeit, 118
erweiterte Zahlengerade, 46
Euler, Leonhard (1707–1783), 83
Euler-MacLaurinsche Summationsformel, 322
Euler-Mascheronische Konstante, 242
Eulersche Beta-Funktion, 254
Eulersche Formel, 145
Eulersche Zahl, 83
Exponentialfunktion zur Basis a , 127
Exponentialreihe, 83
Exponentialreihe im Komplexen, 141

Fakultät, 6
fast alle, 37
Fermat, Pierre de (1601–1665), 13
Fermat-Zahlen, 13
Fibonacci (Leonardo von Pisa) (1180?–1250?), 35
Fibonacci-Zahlen, 35, 47, 68
Fixpunktsatz, 192
Fließkomma, 58
floating point, 58
floor, 32
Folge, 35
Fourier, Joseph (1768–1830), 306
Fourier-Koeffizient, 306
Fourier-Reihe, 306
Freitag, der Dreizehnte, 16
Fresnel, Augustin Jean (1788–1827), 318
Fresnelsche Integrale, 318
Fundamental-Folge, 49
Fundamentalsatz der Differential- und Integralrechnung, 218
Funktionalgleichung der Exponentialfunktion, 87, 142
Funktionalgleichung des Logarithmus, 126

Gamma-Funktion, 242
Gauß, Carl Friedrich (1777–1855), 2, 16
Gauß-Klammer, 32
Gauß'sche Zahlenebene, 137
geometrische Reihe, 11, 44
geometrisches Mittel, 67

- gerade Funktion, 176
 gleichmäßig konvergent, 255
 gleichmäßig stetig, 119
 Gleitpunkt, 58
Gödel, Kurt (1906–1978), 101
 goldener Schnitt, 69
 Graph einer Funktion, 103
*Gregor XIII., Papst (Ugo Buoncom-
 pagni)* (1502–1585), 16
 Gregorianischer Kalender, 16
 Grenzwert, 36
 Grenzwerte bei Funktionen, 106
Hadamard, Jacques S. (1865–1963),
 277
 Hadamardsche Formel, 277
 halboffenes Intervall, 90
 harmonische Reihe, 71
 Häufungspunkt, 56, 94
 Hauptzweige der Arcus-Funktionen,
 156
 Hexadezimalsystem, 60
Hölder, Otto (1859–1937), 185
 Höldersche Ungleichung, 185, 215
*Hospital, Guillaume-François-Antoine
 de l'H.* (1661–1704), 187
 Hospitalsche Regeln, 187
 IEEE-Standard, 59
 imaginäre Einheit, 137
 Imaginärteil, 137
 Indexverschiebung, 10
 Induktion, vollständige, 1
 Induktions-Axiom, 29
 Infimum, 96
 Integral-Sinus, 238
 Integral-Vergleichskriterium, 240
 Intervall-Halbierungsmethode, 113
 Intervallschachtelungs-Prinzip, 50
 Inverses, 19
 irrationale Zahlen, 94
Kepler, Johannes (1571–1630), 234
 Keplersche Fassregel, 234
 Kettenbruch, 68
 Kettenregel, 172
 Kommutativgesetz, 17, 19
 kompaktes Intervall, 117
 komplexe Konjugation, 137
 komplexe Zahlen, 136
 Konjugation, komplexe, 137
 konkav, 182
 Kontinuums-Hypothese, 101
 konvergent, 36
 konvergent, uneigentlich, 45
 Konvergenzradius, 261
 konvex, 182
 konvex, logarithmisch, 243
Kronecker, Leopold (1823–1891), 58
Lagrange, Joseph Louis (1736–1813),
 280
 Lagrangesches Restglied, 280
Landau, Edmund (1877–1938), 131
 Landau-Symbole, 131
 leere Summe, 2
 leeres Produkt, 6
Legrende, Adrien Marie (1752–1833),
 190
 Legendre-Polynome, 232
 Legendresche Differentialgleichung,
 190
 Legrende-Polynom, 190
Leibniz, Gottfried Wilhelm (1646–1716),
 72
 Leibniz'sche Reihe, 74, 275, 289
 Leibniz'sches Konvergenz-Kriterium,
 72

- Limes, 36
limes inferior, 98
limes superior, 98
Lipschitz, Rudolf (1832–1903), 122
Lipschitz-stetig, 122
logarithmisch konvex, 243
Logarithmus, 126
Logarithmus zur Basis a , 133
Logarithmus-Reihe, 285
lokales Maximum, 177
lokales Minimum, 177
- Machin, John* (1685–1751), 289
Machinsche Formel, 289
MacLaurin, Colin (1698–1746), 322
Euler-MacLaurinsche Summationsformel, 322
Majoranten-Kriterium, 75
Mantisse, 58
Mascheroni, Lorenzo (1750–1800), 242
Euler-Mascheronische Konstante, 242
Maximum, 26, 98
lokales, 177
Mersenne, Marin (1588–1648), 13
Mersennesche Primzahlen, 13
Minimum, 26, 98
lokales, 177
Minkowski, Hermann (1864–1909), 185
Minkowskische Ungleichung, 185, 215
Mittelwertsatz, 178
Mittelwertsatz der Integralrechnung, 210
Mittelwertsatz, verallgemeinerter, 191
Møllerup, Johannes (1872–1937), 244
Satz von Bohr/Møllerup, 244
monoton wachsend, fallend, 57
monotone Funktion, 124
natürliche Zahlen, 28
natürlicher Logarithmus, 126
Nebenzweige der Arcus-Funktionen, 157
Negatives, 17
Newton, Isaac (1643–1727), 196
Newtonsches Verfahren, 196
Norm, p -Norm, 184
Nullelement, 17
Nullfolge, 36
Oberintegral, 205
offenes Intervall, 90
Partialbruchzerlegung, 221
Partialsumme, 43
Partielle Integration, 224
Pascal, Blaise (1623–1662), 10
PASCAL-Programme, 93
Pascalsches Dreieck, 10
Peano, Guiseppe (1858–1932), 28
Peano-Axiome, 28
periodische Funktion, 304
periodischer b -adischer Bruch, 60
periodischer Dezimalbruch, 45
 π , 149
Planck, Max (1858–1947), 194
Plancksche Strahlungsfunktion, 194
Polarkoordinaten, 157
Potenz, 22
Potenzreihe, 260
primitive Funktion, 218
Primzahl, 4
Produkt, unendliches, 47
Produktregel, 169

- punktweise konvergent, 255
 quadratische Konvergenz, 65, 198
 quadratisches Mittel, 309
 Quadratwurzel, 62
 Quotienten-Kriterium, 76
 Quotientenregel, 169

Raabe, Josef Ludwig (1801–1859), 81
 Raabesches Konvergenz-Kriterium, 81
 Realteil, 137
 Reihen, unendliche, 43
 relatives Extremum, 177
Riemann, Bernhard (1826–1866), 205
 Riemann-integrierbar, 205
 Riemannsche Summe, 211
 Riemannsche Zetafunktion, 241
Rolle, Michel (1652–1719), 178
 Satz von Rolle, 178

 Schachbrett, 12
 Schnittaxiom, Dedekindsches, 102
 Schranke, 95
Schwarz, Hermann Amandus (1843–1921), 185
 Cauchy-Schwarz'sche Ungleichung, 185
 Secans, 154
 Sexagesimalsystem, 53
Simpson, Thomas (1710–1761), 234
 Simpsonsche Regel, 234
 Sinus, 145
 Sinus hyperbolicus, 111
 Sinus integralis, 238
 Sinus-Produkt, 271
 Stammfunktion, 218

Staudt, Karl Georg Christian von, (1798–1867), 299
 stetig, 109
 stetig differenzierbar, 174
 Stetigkeitsmodul, 122
Stirling, James (1692–1770), 248
 Stirlingsche Formel, 248
 streng monoton wachsend, fallend, 57
 strenges lokales Extremum, 177
 Substitutionsregel, 220
 Summenzeichen, 2
 Supremum, 96
 Supremumsnorm, 258

 Tangens, 153
Taylor, Brook (1685–1731), 279
 Taylor-Polynom, 281
 Taylor-Reihe, 282
 Taylorsche Formel, 279
 Teilfolge, 55
 Teleskop-Summe, 43
 transzendent, 100
 Trapez-Regel, 229
 Treppenfunktion, 105
 trigonometrische Funktionen, 145
 trigonometrisches Polynom, 304
Tschebyscheff, Pafnutij L. (1821–1894), 161
 Tschebyscheff-Polynom, 161
Turing, Alan Mathison (1912–1954), 93
 Turing-Maschine, 93

 überabzählbar, 91
 Umkehrfunktion, 124
 Umordnung von Reihen, 78
 unbestimmtes Integral, 217
 uneigentlich konvergent, 45

- uneigentliches Integral, 235
- uneigentliches Intervall, 90
- unendlich, 46
- unendliche geometrische Reihe, 44
- unendliche Reihen, 43
- unendliches Produkt, 47
- ungerade Funktion, 176
- Unterintegral, 205
- unterliegende Menge einer Folge, 91

- vollständige Induktion, 1
- Vollständigkeits-Axiom, 50
- Vollständigkeitsrelation, 310

- Wallis, John* (1616–1703), 227
- Wallis'sches Produkt, 227
- Weierstraß, Karl* (1815–1897), 55
 - Konvergenzkriterium von Weierstraß, 259
 - Satz von Bolzano-Weierstraß, 55
- Weierstraß'scher Approximationssatz, 295
- Wurzeln, 62, 125

- Zahlenebene, Gauß'sche, 137
- Zahlengerade, 26
- Zahlengerade, erweiterte, 46
- Zetafunktion, Riemannsche, 241
- Zwischenwertsatz, 113

Symbolverzeichnis

$\mathbb{N} = \{0, 1, 2, 3, \dots\}$ = Menge der natürlichen Zahlen

$\mathbb{Z} = \{0, \pm 1, \pm 2, \dots\}$ = Menge der ganzen Zahlen

$\mathbb{Q} = \{\frac{p}{q}: p, q \in \mathbb{Z}, q \neq 0\}$ = Körper der rationalen Zahlen

$\mathbb{R}$ = Körper der reellen Zahlen

$\mathbb{R}^* =$ Menge der reellen Zahlen $\neq 0$

$\mathbb{R}_+ =$ Menge der reellen Zahlen ≥ 0

$\mathbb{R}_+^* =$ Menge der reellen Zahlen > 0

$\mathbb{C}$ = Körper der komplexen Zahlen, 136

$\mathbb{F}_2$ = Körper mit zwei Elementen, 23

$[a, b], [a, b[,]a, b],]a, b[$ Intervalle, 90

$\lfloor x \rfloor = \text{floor}(x)$ = größte ganze Zahl $\leq x$, 32

$\lceil x \rceil = \text{ceil}(x)$ = kleinste ganze Zahl $\geq x$, 32

$[x]$ Gauß-Klammer, alte Bezeichnung für $\lfloor x \rfloor$

$|x|$ Betrag einer reellen oder komplexen Zahl, 29, 138

$\|x\|_p$ p -Norm für Vektoren, 184

$\|f\|_p$ p -Norm für Funktionen, 215

$\|f\|_K$ Supremumsnorm, 258

f'_+, f'_- rechtsseitige (linksseitige) Ableitung, 167

$f|_A$ Beschränkung einer Abbildung $f: X \rightarrow Y$
auf eine Teilmenge $A \subset X$

$a_n \sim b_n$ asymptotische Gleichheit von Folgen, 248

Die üblichen Bezeichnungen aus der Mengenlehre werden als bekannt vorausgesetzt, siehe etwa [Fi], Abschnitt 1.1. Insbesondere ist bei der Teilmengen-Relation $A \subset X$ die Gleichheit $A = X$ zugelassen.