

Kapitel 1

Grundbegriffe

Zu Beginn dieses Kapitels erklären wir die grundlegenden Begriffe der Algebra, die an den verschiedensten Stellen der Mathematik auftauchen. Wie intensiv man dies bei der ersten Lektüre studieren soll, ist eine Frage des Geschmacks. Wer gar keinen Appetit darauf verspürt, kann sofort bis nach 1.4 weiterblättern, wo die wirkliche lineare Algebra, nämlich die Theorie der Vektorräume, beginnt, und das, was er von den Grundbegriffen später benötigt, bei Bedarf nachschlagen.

1.1 Mengen und Abbildungen

Wir wollen hier nicht auf die schwierige Frage eingehen, wie der Begriff „Menge“ erklärt werden kann; das ist Gegenstand der mathematischen Grundlagenforschung. Für die Zwecke der linearen Algebra genügt es, einige elementare Regeln und Bezeichnungen zu erläutern. Wer an einer fundierten Darstellung interessiert ist, möge zum Beispiel [F-P] konsultieren.

1.1.1. Die *endlichen Mengen* kann man im Prinzip durch eine vollständige Liste ihrer Elemente angeben. Man schreibt dafür

$$X := \{x_1, x_2, \dots, x_n\},$$

wobei der Doppelpunkt neben dem Gleichheitszeichen bedeutet, daß das links stehende Symbol X durch den rechts stehenden Ausdruck *definiert* wird. Die x_i heißen *Elemente* von X , in Zeichen $x_i \in X$. Man beachte, daß die Elemente x_i nicht notwendig verschieden sein müssen, und daß die Reihenfolge gleichgültig ist. Man nennt die Elemente $x_1, \dots, x_n \in X$ *paarweise verschieden*, wenn

$x_i \neq x_j$ für $i \neq j$. In diesem Fall ist n die *Anzahl der Elemente* von X .

Die *leere Menge* $\emptyset$ ist dadurch gekennzeichnet, daß sie kein Element enthält.

Eine Menge X' heißt *Teilmenge* von X , in Zeichen $X' \subset X$, wenn aus $x \in X'$ immer $x \in X$ folgt. Es gilt

$$X = Y \Leftrightarrow X \subset Y \text{ und } Y \subset X.$$

Die einfachste *unendliche Menge* ist die Menge

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

der *natürlichen Zahlen*. Man kann sie charakterisieren durch die PEANO-Axiome (vgl. [Z]). Diese enthalten das *Prinzip der vollständigen Induktion*:

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

a) $0 \in M$,

b) $n \in M \Rightarrow n + 1 \in M$.

Dann ist $M = \mathbb{N}$.

Mancher Leser wird sich an dieser Stelle zum ersten Mal – aber sicher insgesamt nicht zum letzten Mal – darüber wundern, daß die Bezeichnungen in der Mathematik nicht einheitlich sind. So wird die Null manchmal nicht als natürliche Zahl angesehen. Es gibt Versuche, hier durch DIN-Normen Ordnung zu schaffen (vgl. [DIN]), aber viele Mathematiker lieben mehr ihre Freiheit, als Normblätter.

Durch *Erweiterungen von Zahlbereichen* erhält man ausgehend von $\mathbb{N}$ die *ganzen Zahlen*

$$\mathbb{Z} = \{0, +1, -1, +2, -2, \dots\},$$

die *rationalen Zahlen*

$$\mathbb{Q} = \left\{ \frac{p}{q} : p, q \in \mathbb{Z}, q \neq 0 \right\},$$

und etwa als Dezimalbrüche oder Cauchy-Folgen rationaler Zahlen die *reellen Zahlen* $\mathbb{R}$. Es ist

$$\mathbb{N} \subset \mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R} \subset \mathbb{C},$$

wobei die in gewisser Weise abschließende Erweiterung zu den komplexen Zahlen $\mathbb{C}$ in 1.3.3 und 1.3.9 behandelt wird. Einem Leser, der mehr über den Aufbau der „Zahlen“ wissen möchte, sei [Z] empfohlen.

1.1.2. In der linearen Algebra werden wir mit solchen Mengen auskommen, die sich aus den in 1.1.1 angegebenen Mengen mit Hilfe einiger elementarer Operationen ableiten lassen.

Aus einer gegebenen Menge X kann man *Teilmengen* auswählen, die durch gewisse Eigenschaften der Elemente charakterisiert sind, in Zeichen

$$X' := \{x \in X : x \text{ hat die Eigenschaft } E\},$$

zum Beispiel $X' := \{n \in \mathbb{N} : n \text{ ist gerade}\}$.

Sind $X_1, \dots, X_n$ Mengen, so hat man eine *Vereinigung*

$$X_1 \cup \dots \cup X_n := \{x : \text{es gibt ein } i \in \{1, \dots, n\} \text{ mit } x \in X_i\}$$

und den *Durchschnitt*

$$X_1 \cap \dots \cap X_n := \{x : x \in X_i \text{ für alle } i \in \{1, \dots, n\}\}.$$

An einigen Stellen wird es nötig sein, Vereinigungen und Durchschnitte von mehr als endlich vielen Mengen zu betrachten. Dazu verwendet man eine Menge I , die *Indexmenge* heißt, so daß für jedes $i \in I$ eine Menge X_i gegeben ist. Dann sind *Vereinigung* und *Durchschnitt* erklärt durch

$$\bigcup_{i \in I} X_i := \{x: \text{ es gibt ein } i \in I \text{ mit } x \in X_i\},$$

$$\bigcap_{i \in I} X_i := \{x: x \in X_i \text{ für alle } i \in I\}.$$

Ist etwa $I = \mathbb{N}$ und $X_i := [-i, i] \subset \mathbb{R}$ für jedes $i \in \mathbb{N}$ ein Intervall, so ist

$$\bigcup_{i \in \mathbb{N}} X_i = \mathbb{R} \quad \text{und} \quad \bigcap_{i \in \mathbb{N}} X_i = \{0\}.$$

Ist $X' \subset X$ eine Teilmenge, so nennt man

$$X \setminus X' := \{x \in X: x \notin X'\}$$

die *Differenzmenge* (oder das *Komplement*).

1.1.3. Zur Beschreibung von Beziehungen zwischen verschiedenen Mengen verwendet man „Abbildungen“. Sind X und Y Mengen, so versteht man unter einer *Abbildung* von X nach Y eine Vorschrift f , die jedem $x \in X$ eindeutig ein $f(x) \in Y$ zuordnet. Man schreibt dafür

$$f: X \rightarrow Y, \quad x \mapsto f(x).$$

Man beachte den Unterschied zwischen den beiden Pfeilen: „ $\rightarrow$ “ steht zwischen den Mengen und „ $\mapsto$ “ zwischen den Elementen.

Zwei Abbildungen $f: X \rightarrow Y$ und $g: X \rightarrow Y$ heißen *gleich*, in Zeichen $f = g$, wenn $f(x) = g(x)$ für alle $x \in X$. Mit

$$\text{Abb}(X, Y)$$

bezeichnen wir die Menge aller Abbildungen von X nach Y .

Ein Problem bei dieser Definition einer Abbildung ist, daß nicht präzisiert ist, in welcher Form die Abbildungsvorschrift gegeben sein soll (genauso wenig war festgelegt worden, in welcher Form die charakterisierende Eigenschaft einer Teilmenge vorliegen soll). Besonders einfach zu beschreiben sind Abbildungen, bei denen $f(x)$ durch eine explizite Formel angegeben werden kann, etwa im Fall $X = Y = \mathbb{R}$ durch

$$f(x) = ax, \quad f(x) = x^2, \quad f(x) = \sqrt{x}.$$

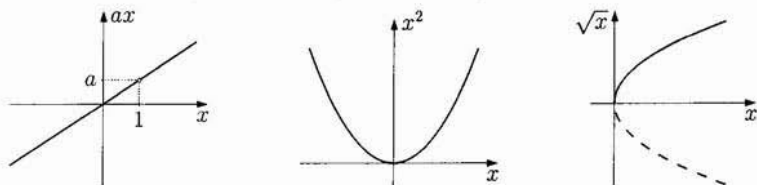


Bild 1.1

Die letzte Vorschrift ist schon problematisch, weil es für positive reelle Zahlen zwei und für negative keine Quadratwurzel gibt. Eine Abbildung im Sinn der Definition liegt nur vor, wenn man negative x ausschließt und für positive x eine Wurzel (etwa die positive) auswählt. Dann hat man eine Abbildung

$$f: \mathbb{R}_+ \rightarrow \mathbb{R}_+, \quad x \mapsto +\sqrt{x},$$

wobei $\mathbb{R}_+ := \{x \in \mathbb{R}: x \geq 0\}$.

Mit einer Abbildung kann man nicht nur Elemente, sondern auch Teilmengen bewegen. Sei also $f: X \rightarrow Y$ und $M \subset X$, $N \subset Y$. Dann heißt

$$f(M) := \{y \in Y: \text{es gibt ein } x \in M \text{ mit } y = f(x)\} \subset Y$$

das *Bild* von M (in Y), insbesondere $f(X)$ das Bild von X in Y , und

$$f^{-1}(N) := \{x \in X: f(x) \in N\} \subset X$$

das *Urbild* von N in X . Man beachte, daß für eine einelementige Menge $N = \{y\}$ das Urbild

$$f^{-1}(y) := f^{-1}(\{y\}) \subset X$$

eine Teilmenge ist, die aus mehreren Elementen bestehen, aber auch leer sein kann (etwa bei $f(x) = x^2$). Daher ist f^{-1} im allgemeinen keine Abbildung von Y nach X im Sinn der Definition.

Noch eine Bezeichnungsweise: Ist $f: X \rightarrow Y$ eine Abbildung und $M \subset X$ eine Teilmenge, so nennt man

$$f|_M: M \rightarrow Y, \quad x \mapsto f(x),$$

die *Beschränkung* von f auf M . Sie unterscheidet sich von f nur durch den eingeschränkten Definitionsbereich. Ist $Y \subset Y'$ eine Teilmenge, so ist es üblich, mit $f: X \rightarrow Y'$ auch die Abbildung mit dem ausgedehnten Bildbereich zu bezeichnen.

1.1.4. Besonders wichtige Eigenschaften von Abbildungen haben eigene Namen. Eine Abbildung $f: X \rightarrow Y$ heißt

injektiv, falls aus $x, x' \in X$ und $f(x) = f(x')$ stets $x = x'$ folgt,

surjektiv, falls $f(X) = Y$,

d.h. falls es zu jedem $y \in Y$ ein $x \in X$ gibt mit $y = f(x)$,

bijektiv, falls f injektiv und surjektiv ist.

Ist f bijektiv, so gibt es zu jedem $y \in Y$ genau ein $x \in X$ mit $f(x) = y$. In diesem Fall kann man also eine *Umkehrabbildung*

$$f^{-1}: Y \rightarrow X, \quad y \mapsto x = f^{-1}(y) \quad \text{mit} \quad y = f(x)$$

erklären. Man beachte, daß das Symbol f^{-1} in verschiedenen Bedeutungen vorkommt:

- bei einer beliebigen Abbildung ist für jede Teilmenge $N \subset Y$ das Urbild $f^{-1}(N) \subset X$ eine Teilmenge,
- ist f bijektiv, so besteht für die einelementige Teilmenge $\{y\} \subset Y$ das Urbild $f^{-1}(\{y\})$ aus einem Element x , in Zeichen

$$f^{-1}(\{y\}) = \{x\},$$

dafür schreibt man dann $f^{-1}(y) = x$.

Durch systematisches Zählen beweist man den folgenden einfachen

Satz. Ist X eine endliche Menge, so sind für eine Abbildung $f: X \rightarrow X$ folgende Bedingungen äquivalent:

- i) f ist injektiv,
- ii) f ist surjektiv,
- iii) f ist bijektiv.

Beweis. X bestehe aus n Elementen, also $X = \{x_1, \dots, x_n\}$, wobei $n \in \mathbb{N}$ ist und die x_i paarweise verschieden sind.

i) $\Rightarrow$ ii): Wir zeigen „nicht surjektiv“ $\Rightarrow$ „nicht injektiv“. Ist $f(X) \neq X$, so besteht $f(X)$ aus $m < n$ Elementen. Nun besagt das sogenannte *Schubladenprinzip von DIRICHLET*: Verteilt man n Objekte irgendwie in m Schubladen, wobei $m < n$, so gibt es mindestens eine Schublade, in der mehr als ein Objekt liegt. Also kann f nicht injektiv sein.

ii) $\Rightarrow$ i): Ist f nicht injektiv, so gibt es x_i, x_j mit $x_i \neq x_j$, aber $f(x_i) = f(x_j)$. Dann kann $f(X)$ höchstens $n - 1$ Elemente enthalten, also ist f nicht surjektiv.

Wegen ii) $\Rightarrow$ i) folgt auch ii) $\Rightarrow$ iii), iii) $\Rightarrow$ i) ist klar. Damit sind alle möglichen Implikationen bewiesen. $\square$

1.1.5. Sind X, Y, Z Mengen und $f: X \rightarrow Y$ sowie $g: Y \rightarrow Z$ Abbildungen, so heißt die Abbildung

$$g \circ f: X \rightarrow Z, \quad x \mapsto g(f(x)) =: (g \circ f)(x)$$

die *Komposition* (oder *Hintereinanderschaltung*) von f und g (man sagt g *kompontiert mit f* für $g \circ f$). Man beachte dabei, daß die zuerst angewandte Abbildung rechts steht, im Gegensatz zum Diagramm

$$X \xrightarrow{f} Y \xrightarrow{g} Z.$$

Daher könnte man das Diagramm besser umgekehrt schreiben:

$$Z \xleftarrow{g} Y \xleftarrow{f} X.$$

Bemerkung. Die Komposition von Abbildungen ist assoziativ, d.h. für Abbildungen $f: X \rightarrow Y$, $g: Y \rightarrow Z$ und $h: Z \rightarrow W$ ist

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

Beweis. Ist $x \in X$, so ist

$$\begin{aligned} ((h \circ g) \circ f)(x) &= (h \circ g)(f(x)) = h(g(f(x))) \\ &= h((g \circ f)(x)) = (h \circ (g \circ f))(x). \end{aligned} \quad \square$$

Vorsicht! Die Komposition von Abbildungen ist i.a. *nicht* kommutativ. Ist

$$\begin{aligned} f: \mathbb{R} &\rightarrow \mathbb{R}, & x &\mapsto x + 1, \\ g: \mathbb{R} &\rightarrow \mathbb{R}, & x &\mapsto x^2, \end{aligned}$$

so ist $(f \circ g)(x) = x^2 + 1$ und $(g \circ f)(x) = (x + 1)^2$, also $f \circ g \neq g \circ f$.

Um eine etwas andersartige Charakterisierung von Injektivität und Surjektivität zu erhalten, erklären wir für jede Menge X die identische Abbildung

$$\text{id}_X: X \rightarrow X, \quad x \mapsto x.$$

Lemma. Sei $f: X \rightarrow Y$ eine Abbildung zwischen den nicht leeren Mengen X und Y . Dann gilt:

- 1) f ist genau dann injektiv, wenn es eine Abbildung $g: Y \rightarrow X$ gibt, so daß $g \circ f = \text{id}_X$.
- 2) f ist genau dann surjektiv, wenn es eine Abbildung $g: Y \rightarrow X$ gibt, so daß $f \circ g = \text{id}_Y$.
- 3) f ist genau dann bijektiv, wenn es eine Abbildung $g: Y \rightarrow X$ gibt, so daß $g \circ f = \text{id}_X$ und $f \circ g = \text{id}_Y$. In diesem Fall ist $g = f^{-1}$.

Beweis. 1) Sei f injektiv. Dann gibt es zu jedem $y \in f(X)$ genau ein $x \in X$ mit $f(x) = y$, und wir definieren $g(y) := x$. Ist $x_0 \in X$ beliebig, so definieren wir weiter $g(y) = x_0$ für alle $y \in Y \setminus f(X)$. Das ergibt eine Abbildung $g: Y \rightarrow X$ mit $g \circ f = \text{id}_X$.

Ist umgekehrt $g: Y \rightarrow X$ mit $g \circ f = \text{id}_X$ gegeben, und ist $f(x) = f(x')$ für $x, x' \in X$, so ist $x = \text{id}_X(x) = g(f(x)) = g(f(x')) = \text{id}_X(x') = x'$. Also ist f injektiv.

2) Sei f surjektiv. Zu jedem $y \in Y$ wählen wir ein festes $x \in X$ mit $f(x) = y$ und setzen $g(y) := x$. Die so erklärte Abbildung $g: Y \rightarrow X$ hat die Eigenschaft $f \circ g = \text{id}_Y$.

Ist umgekehrt $g: Y \rightarrow X$ mit $f \circ g = \text{id}_Y$ gegeben, und ist $y \in Y$, so ist $y = f(g(y))$, also Bild von $g(y)$, und f ist surjektiv.

3) Ist f bijektiv, so erfüllt $g := f^{-1}$ die beiden Gleichungen. Ist umgekehrt $g: Y \rightarrow X$ mit $g \circ f = \text{id}_X$ und $f \circ g = \text{id}_Y$ gegeben, so ist f nach 1) und 2) bijektiv, und es gilt $g = f^{-1}$. $\square$

1.1.6. Schon bei der Einführung des Raumes $\mathbb{R}^n$ hatten wir ein „direktes Produkt“ betrachtet. Sind allgemeiner $X_1, \dots, X_n$ Mengen, so betrachten wir die *geordneten n -Tupel*

$$x = (x_1, \dots, x_n) \quad \text{mit} \quad x_i \in X_1, \dots, x_n \in X_n.$$

Genauer kann man x als Abbildung

$$x: \{1, \dots, n\} \rightarrow X_1 \cup \dots \cup X_n \quad \text{mit} \quad x(i) \in X_i$$

ansetzen (man nennt x auch *Auswahlfunktion*), und zur Vereinfachung der Schreibweise $x_i := x(i)$ und $x := (x_1, \dots, x_n)$ setzen. Im Sinne der Gleichheit von Abbildungen gilt dann

$$(x_1, \dots, x_n) = (x'_1, \dots, x'_n) \Leftrightarrow x_1 = x'_1, \dots, x_n = x'_n.$$

Nun erklären wir das *direkte Produkt*

$$X_1 \times \dots \times X_n := \{(x_1, \dots, x_n) : x_i \in X_i\}$$

als Menge der geordneten n -Tupel. Offensichtlich ist $X_1 \times \dots \times X_n \neq \emptyset$, wenn $X_i \neq \emptyset$ für alle $i \in \{1, \dots, n\}$.

Für jedes i hat man eine *Projektion* auf die i -te *Komponente*

$$\pi_i: X_1 \times \dots \times X_n \rightarrow X_i, \quad (x_1, \dots, x_i, \dots, x_n) \mapsto x_i.$$

Ist speziell $X_1 = \dots = X_n = X$, so schreibt man

$$X^n = X \times \dots \times X.$$

Ein Element von X^n ist also eine Abbildung $\{1, \dots, n\} \rightarrow X$. Ist allgemeiner I irgendeine Menge (man nennt sie *Indexmenge*), so nennt man eine Abbildung

$$\varphi: I \rightarrow X, \quad i \mapsto x_i = \varphi(i),$$

eine *Familie* von Elementen in X . Man beachte den Unterschied zu einer Teilmenge $X' \subset X$, die man mit $\varphi(I)$ vergleichen kann: die Elemente $i \in I$ kann man als Etiketten (oder noch näher am Familienleben als Pflichten) ansehen, die unter den Elementen von X verteilt werden. Jedes Etikett i hat einen eindeutigen Empfänger x_i , die Elemente von $X' = \varphi(I)$ erhalten mindestens ein Etikett,

und es ist möglich, daß manche $x \in X'$ mehrere Etiketten erhalten (was zur oft gehörten Klage „immer ich“ führt).

Zur Abkürzung bezeichnet man eine Familie $I \rightarrow X$ oft mit $(x_i)_{i \in I}$.

Für die Indexmenge $I = \mathbb{N}$ der natürlichen Zahlen nennt man $(x_i)_{i \in \mathbb{N}}$ eine *Folge*, das ist ein grundlegender Begriff der Analysis.

Vorsicht! Die Frage der Existenz der oben betrachteten Auswahlfunktionen ist nicht unproblematisch. Das führt zum *Auswahlaxiom*, vgl. [F-P], das auch im Beweis von Lemma 1.1.5 schon stillschweigend verwendet wurde.

1.1.7. Für eine Abbildung $f: X \rightarrow Y$ nennt man die Menge

$$\Gamma_f := \{(x, f(x)) \in X \times Y\}$$

den *Graphen* von f . Damit kann man oft eine Abbildung durch eine Skizze veranschaulichen.

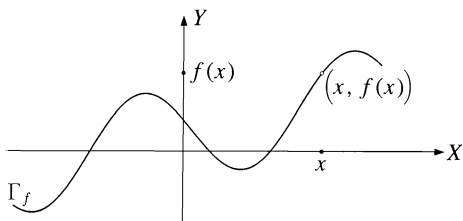


Bild 1.2

Nach der Definition einer Abbildung ist die Einschränkung der Projektion auf X

$$\pi: \Gamma_f \rightarrow X$$

bijektiv. Daraus folgt, daß das „Funktionengebirge“ keine „Überhänge“ hat, wie im folgenden Bild:



Bild 1.3 Gebirge mit Überhängen: In der Fränkischen Schweiz

1.1.8. Das direkte Produkt ist auch nützlich, um Beziehungen (oder Relationen) zwischen je zwei Elementen $x, y \in X$ zu studieren. Man schreibt dafür allgemein $x \sim y$.

Beispiele. a) $X =$ Menge der Menschen, $x \sim y :\Leftrightarrow x$ kennt y .

b) $X = \mathbb{R}$, $x \sim y :\Leftrightarrow x \leq y$.

c) $X = \mathbb{R}^n$, $(x_1, \dots, x_n) \sim (y_1, \dots, y_n) :\Leftrightarrow x_1^2 + \dots + x_n^2 = y_1^2 + \dots + y_n^2$.

d) $X = \mathbb{Z}$, $0 \neq m \in \mathbb{N}$, $x \sim y :\Leftrightarrow y - x$ durch m teilbar.

Eine Relation ist beschrieben durch ihren *Graphen* $R \subset X \times X$, wobei

$$(x, y) \in R \Leftrightarrow x \sim y. \quad (*)$$

Man kann also eine *Relation* auf X definieren als Teilmenge $R \subset X \times X$, und das Zeichen $\sim$ durch $(*)$ erklären.

Definition. Eine Relation $\sim$ auf X heißt *Äquivalenzrelation*, wenn für beliebige $x, y, z \in X$ gilt:

A1 $x \sim x$, (reflexiv)

A2 $x \sim y \Rightarrow y \sim x$, (symmetrisch)

A3 $x \sim y$ und $y \sim z \Rightarrow x \sim z$. (transitiv)

In diesem Fall sagt man x ist *äquivalent* zu y für $x \sim y$.

Der Leser möge zur Übung die obigen Beispiele auf diese Eigenschaften überprüfen. Vor allem die Reflexivität in Beispiel a) ist etwas Nachdenken wert.

Hat man auf einer Menge eine Äquivalenzrelation eingeführt, so kann man – wie wir sehen werden – zu einer neuen Menge übergehen, in der äquivalente Elemente der ursprünglichen Menge zu „Repräsentanten“ desselben neuen Elementes werden. Dabei wird – mit den Worten von HERMANN WEYL – *alles im Sinne des eingenommenen Standpunktes Unwesentliche an den untersuchten Objekten abgestreift*. Übersetzt ins Mengen-Latein, liest sich das wie folgt:

Ist eine Äquivalenzrelation auf einer Menge X gegeben, so heißt eine Teilmenge $A \subset X$ *Äquivalenzklasse* (bezüglich R), falls gilt:

1. $A \neq \emptyset$.

2. $x, y \in A \Rightarrow x \sim y$.

3. $x \in A, y \in X, x \sim y \Rightarrow y \in A$.

Man überlege sich, wie die Äquivalenzklassen in obigen Beispielen c) und d) aussehen.

Bemerkung. Ist R eine Äquivalenzrelation auf einer Menge X , so gehört jedes Element $a \in X$ zu genau einer Äquivalenzklasse. Insbesondere gilt für zwei beliebige Äquivalenzklassen A, A' entweder $A = A'$ oder $A \cap A' = \emptyset$.

Beweis. Für ein fest gegebenes $a \in X$ definieren wir

$$A := \{x \in X : x \sim a\}.$$

Wir zeigen, daß A eine Äquivalenzklasse ist, die a enthält. Wegen $a \sim a$ ist $a \in A$, und es folgt $A \neq \emptyset$. Sind $x, y \in A$, so ist $x \sim a$ und $y \sim a$, also $x \sim y$ nach A2 und A3.

Ist $x \in A$, $y \in X$ und $x \sim y$, so ist $x \sim a$, also nach A2 und A3 auch $y \sim a$ und daher $y \in A$. Damit ist gezeigt, daß a in mindestens einer Äquivalenzklasse enthalten ist.

Es bleibt zu zeigen, daß zwei Äquivalenzklassen A und A' entweder gleich oder disjunkt sind. Angenommen, es ist $A \cap A' \neq \emptyset$ und $a \in A \cap A'$. Ist $x \in A$, so ist $x \sim a$, und wegen $a \in A'$ folgt auch $x \in A'$. Also ist $A \subset A'$. Ebenso beweist man $A' \subset A$, woraus $A = A'$ folgt. $\square$

Jede Äquivalenzrelation R auf einer Menge X liefert also eine *Zerlegung* von X in disjunkte Äquivalenzklassen. Diese Äquivalenzklassen betrachtet man als Elemente einer neuen Menge, die man mit X/R bezeichnet. Man nennt sie die *Quotientenmenge* von X nach der Äquivalenzrelation R . Die *Elemente* von X/R sind also spezielle *Teilmengen* von X . Indem man jedem Element $a \in X$ die Äquivalenzklasse A_a zuordnet, in der es enthalten ist, erhält man eine *kanonische* (d.h. in der gegebenen Situation eindeutig festgelegte) Abbildung

$$X \rightarrow X/R, \quad a \mapsto A_a.$$

Das Urbild eines Elementes $A \in X/R$ ist dabei wieder A , aber aufgefaßt als *Teilmenge* von X .

Jedes $a \in A$ heißt ein *Repräsentant* der Äquivalenzklasse A . Im allgemeinen gibt es keine Möglichkeit, spezielle Repräsentanten besonders auszuzeichnen. Das wäre auch gar nicht im Sinne der durchgeführten Konstruktion.

Als Beispiel aus dem realen Leben kann eine Schule dienen: die Menge der Schüler wird in Klassen eingeteilt, und für manche Fragen, etwa die Gestaltung des Stundenplans, ist nur noch die Menge der Klassen von Bedeutung.

Aufgaben zu 1.1

1. Beweisen Sie die folgenden Rechenregeln für die Operationen mit Mengen:

- $X \cap Y = Y \cap X$, $X \cup Y = Y \cup X$,
- $X \cap (Y \cap Z) = (X \cap Y) \cap Z$, $X \cup (Y \cup Z) = (X \cup Y) \cup Z$,

- c) $X \cap (Y \cup Z) = (X \cap Y) \cup (X \cap Z)$, $X \cup (Y \cap Z) = (X \cup Y) \cap (X \cup Z)$,
 d) $X \setminus (M_1 \cap M_2) = (X \setminus M_1) \cup (X \setminus M_2)$, $X \setminus (M_1 \cup M_2) = (X \setminus M_1) \cap (X \setminus M_2)$.

2. Sei $f: X \rightarrow Y$ eine Abbildung. Zeigen Sie:

- a) Ist $M_1 \subset M_2 \subset X$, so folgt $f(M_1) \subset f(M_2)$.
 Ist $N_1 \subset N_2 \subset Y$, so folgt $f^{-1}(N_1) \subset f^{-1}(N_2)$.
 b) $M \subset f^{-1}(f(M))$ für $M \subset X$, $f(f^{-1}(N)) \subset N$ für $N \subset Y$.
 c) $f^{-1}(Y \setminus N) = X \setminus f^{-1}(N)$ für $N \subset Y$.
 d) Für $M_1, M_2 \subset X$ und $N_1, N_2 \subset Y$ gilt:

$$f^{-1}(N_1 \cap N_2) = f^{-1}(N_1) \cap f^{-1}(N_2), \quad f^{-1}(N_1 \cup N_2) = f^{-1}(N_1) \cup f^{-1}(N_2),$$

$$f(M_1 \cup M_2) = f(M_1) \cup f(M_2), \quad f(M_1 \cap M_2) \subset f(M_1) \cap f(M_2).$$

Finden Sie ein Beispiel, in dem $f(M_1 \cap M_2) \neq f(M_1) \cap f(M_2)$ gilt!

3. Seien $f: X \rightarrow Y$, $g: Y \rightarrow Z$ Abbildungen und $g \circ f: X \rightarrow Z$ die Komposition von f und g . Dann gilt:

- a) Sind f und g injektiv (surjektiv), so ist auch $g \circ f$ injektiv (surjektiv).
 b) Ist $g \circ f$ injektiv (surjektiv), so ist auch f injektiv (g surjektiv).

4. Untersuchen Sie die folgenden Abbildungen auf Injektivität und Surjektivität:

- a) $f_1: \mathbb{R}^2 \rightarrow \mathbb{R}$, $(x, y) \mapsto x + y$, b) $f_2: \mathbb{R}^2 \rightarrow \mathbb{R}$, $(x, y) \mapsto x^2 + y^2 - 1$,
 c) $f_3: \mathbb{R}^2 \rightarrow \mathbb{R}^2$, $(x, y) \mapsto (x + 2y, 2x - y)$,

5. Zwei Mengen X und Y heißen gleichmächtig genau dann, wenn es eine bijektive Abbildung $f: X \rightarrow Y$ gibt. Eine Menge X heißt abzählbar unendlich, falls X und $\mathbb{N}$ gleichmächtig sind.

- a) Zeigen Sie, dass $\mathbb{Z}$ und $\mathbb{Q}$ abzählbar unendlich sind.
 b) Zeigen Sie, dass $\mathbb{R}$ nicht abzählbar unendlich ist.
 c) Für eine nichtleere Menge M sei $\text{Abb}(M, \{0, 1\})$ die Menge aller Abbildungen von M nach $\{0, 1\}$. Zeigen Sie, dass M und $\text{Abb}(M, \{0, 1\})$ nicht gleichmächtig sind.

6. Ein Konferenzhotel für Mathematiker hat genau $\mathbb{N}$ Betten. Das Hotel ist bereits voll belegt, aber die Mathematiker lassen sich nach Belieben innerhalb des Hotels umquartieren. Das Hotel soll aus wirtschaftlichen Gründen stets voll belegt sein, und wenn möglich, sollen alle neu ankommenden Gäste untergebracht werden. Was macht man in folgenden Fällen?

- a) Ein weiterer Mathematiker trifft ein.
 b) Die Insassen eines Kleinbusses mit n Plätzen suchen Unterkunft.
 c) Ein Großraumbus mit $\mathbb{N}$ Personen kommt an.
 d) n Großraumbusse treffen ein.
 e) $\mathbb{N}$ Großraumbusse fahren vor.

1.2 Gruppen

1.2.1. Unter einer *Verknüpfung* (oder *Komposition*) auf einer Menge G versteht man eine Vorschrift $*$, die zwei gegebenen Elementen $a, b \in G$ ein neues Element $*(a, b) \in G$ zuordnet, d.h. eine Abbildung

$$*: G \times G \rightarrow G, \quad (a, b) \mapsto *(a, b).$$

Wir geben einige Beispiele für solche Vorschriften $*$ und schreiben dabei zur Abkürzung $a * b$ statt $*(a, b)$:

a) Ist X eine Menge und $G = \text{Abb}(X, X)$ die Menge aller Abbildungen

$$f: X \rightarrow X,$$

so ist für $f, g \in G$ nach 1.1.5 auch $f \circ g \in G$, also kann man

$$f * g := f \circ g$$

setzen.

b) In $G = \mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}$ oder $\mathbb{R}_+^*$ hat man Addition und Multiplikation, also kann man für $a, b \in G$

$$a * b := a + b \quad \text{oder} \quad a * b := a \cdot b$$

setzen. Im Gegensatz zu Beispiel a) ist die Reihenfolge von a und b hier egal.

c) In $G = \mathbb{Q}$ oder $\mathbb{R}$ kann man

$$a * b := \frac{1}{2}(a + b)$$

als das *arithmetische Mittel* von a und b erklären.

1.2.2. Verknüpfungen mit besonders guten Eigenschaften haben eigene Namen.

Definition. Eine Menge G zusammen mit einer Verknüpfung $*$ heißt *Gruppe*, wenn folgende Axiome erfüllt sind:

G1 $(a * b) * c = a * (b * c)$ für alle $a, b, c \in G$ (*Assoziativgesetz*).

G2 Es gibt ein $e \in G$ (*neutrales Element* genannt) mit den folgenden Eigenschaften:

a) $e * a = a$ für alle $a \in G$.

b) Zu jedem $a \in G$ gibt es ein $a' \in G$ (*inverses Element von a* genannt) mit $a' * a = e$.

Die Gruppe heißt *abelsch* (oder *kommutativ*), falls außerdem $a * b = b * a$ für alle $a, b \in G$.

Falls das keine Verwirrung stiftet, schreibt man die Verknüpfung zur Vereinfachung meist als Multiplikation, also $a \cdot b$ oder nur ab statt $a * b$.

Ist die Verknüpfung als Addition geschrieben, so setzt man stillschweigend voraus, daß sie kommutativ ist. Das neutrale Element 0 heißt dann *Nullelement*, das Inverse von a wird mit $-a$ bezeichnet und heißt *Negatives*.

Wenn das Assoziativgesetz erfüllt ist, kann man bei mehrfachen Produkten die Klammern weglassen, also schreiben:

$$abc = (ab)c = a(bc).$$

Zunächst wollen wir die Gruppenaxiome bei den Beispielen aus 1.2.1 nachprüfen.

a) In $G = \text{Abb}(X, X)$ ist die Komposition nach 1.1.5 assoziativ, die identische Abbildung id_X erfüllt G2a, aber aus der Existenz eines g mit $g \circ f = \text{id}_X$ folgt, daß f injektiv sein muß. Also ist G2b im allgemeinen nicht erfüllt.

Das kann man aber reparieren. Die Hintereinanderschaltung ist auch eine Verknüpfung in der Teilmenge

$$S(X) := \{f \in \text{Abb}(X, X) : f \text{ bijektiv}\},$$

und $S(X)$ wird damit zu einer Gruppe. Sie heißt die *symmetrische Gruppe* der Menge X . Ist $X = \{1, \dots, n\}$, so schreibt man S_n statt $S(X)$. Jedes $\sigma \in S_n$ heißt *Permutation* der Zahlen $1, \dots, n$, und S_n nennt man auch *Permutationsgruppe*. In der linearen Algebra ist sie wichtig bei der Berechnung von Determinanten, daher verschieben wir alles Weitere hierüber auf Kapitel 3. Insbesondere werden wir dort sehen, daß S_n für $n \geq 3$ nicht abelsch ist.

b) Hier sind nur $\mathbb{Z}$, $\mathbb{Q}$ und $\mathbb{R}$ mit der Addition und $\mathbb{R}_+^*$ mit der Multiplikation Gruppen. Der Leser möge zur Übung nachprüfen, welche Axiome in den anderen Fällen verletzt sind.

c) Das arithmetische Mittel ist nur kommutativ, aber nicht assoziativ, und es gibt kein neutrales Element.

1.2.3. Bei der Aufstellung von Axiomen versucht man, so wenig wie möglich zu fordern und die weiteren Eigenschaften daraus abzuleiten. Insbesondere haben wir weder bei e noch bei a' die Eindeutigkeit postuliert. Derartigen Kleinkram kann man nachträglich beweisen.

Bemerkung. Ist G eine Gruppe, so gilt:

a) Das neutrale Element $e \in G$ ist eindeutig bestimmt und hat auch die Eigenschaft $a \cdot e = a$ für alle $a \in G$.

b) Das inverse Element a' ist für jedes $a \in G$ eindeutig bestimmt und hat auch die Eigenschaft $a \cdot a' = e$ für alle $a \in G$.

c) Da das Inverse zu a nach b) eindeutig bestimmt ist, kann man es mit a^{-1} bezeichnen. Es gilt für $a, b \in G$:

$$a^{-1} \cdot a = a \cdot a^{-1} = e, \quad (a^{-1})^{-1} = a, \quad (ab)^{-1} = b^{-1}a^{-1}.$$

d) Es gelten die folgenden Kürzungsregeln:

$$a \cdot \tilde{x} = a \cdot x \Rightarrow x = \tilde{x} \quad \text{und} \quad y \cdot a = \tilde{y} \cdot a \Rightarrow y = \tilde{y}.$$

Beweis. Wir betrachten ein neutrales e und ein $a \in G$. Zu a' gibt es ein a'' mit $a''a' = e$. Daraus folgt

$$\begin{aligned} aa' &= e(aa') = (a''a')(aa') = a''(a'(aa')) \\ &= a''((a'a)a') = a''(ea') = a''a' = e, \end{aligned}$$

und somit $ae = a(a'a) = (aa')a = ea = a$.

Ist $\tilde{e}$ ein eventuelles anderes neutrales Element, so ist

$$e\tilde{e} = e \quad \text{und} \quad e\tilde{e} = \tilde{e}, \quad \text{also} \quad e = \tilde{e}.$$

Damit ist a) und die erste Gleichung von c) bewiesen.

Ist $\tilde{a}'$ ein eventuelles anderes Inverses, so folgt

$$\tilde{a}' = \tilde{a}'e = \tilde{a}'(aa') = (\tilde{a}'a)a' = ea' = a'$$

unter Verwendung der bereits vorher bewiesenen Eigenschaften. Damit ist auch b) bewiesen.

Aus $aa^{-1} = e$ folgt, daß a inverses Element zu a^{-1} ist, d.h. $(a^{-1})^{-1} = a$. Weiter gilt

$$(b^{-1}a^{-1})(ab) = b^{-1}(a^{-1}(ab)) = b^{-1}((a^{-1}a)b) = b^{-1}(eb) = b^{-1}b = e.$$

Schließlich folgen die Kürzungsregeln durch Multiplikation der jeweils ersten Gleichung von links bzw. rechts mit a^{-1} . $\square$

1.2.4. Auf der Suche nach Beispielen für Gruppen kann es helfen, das Axiom G2 etwas anders zu formulieren. Dazu betrachten wir für ein festes $a \in G$ die Abbildungen

$$\begin{aligned} \tau_a: G &\rightarrow G, \quad x \mapsto x \cdot a, \quad (\text{Rechtstranslation}), \text{ und} \\ {}_a\tau: G &\rightarrow G, \quad x \mapsto a \cdot x, \quad (\text{Linkstranslation}). \end{aligned}$$

Lemma. Ist G eine Gruppe, so sind für jedes $a \in G$ die Abbildungen τ_a und ${}_a\tau$ bijektiv.

Ist umgekehrt G eine Menge mit einer assoziativen Verknüpfung, so folgt G2 aus der Surjektivität der Abbildungen τ_a und ${}_a\tau$ für alle $a \in G$.

Beweis. Die Bijektivität von τ_a und ${}_a\tau$ bedeutet, daß es zu $b \in G$ genau ein $x \in G$ und genau ein $y \in G$ gibt mit

$$x \cdot a = b \quad \text{und} \quad a \cdot y = b,$$

d.h. daß diese Gleichungen mit x und y als Unbestimmten eindeutig lösbar sind. Die Existenz einer Lösung ist klar, denn es genügt,

$$x = b \cdot a^{-1} \quad \text{und} \quad y = a^{-1} \cdot b$$

zu setzen. Sind $\tilde{x}$ und $\tilde{y}$ weitere Lösungen, so gilt

$$x \cdot a = \tilde{x} \cdot a \quad \text{und} \quad a \cdot y = a \cdot \tilde{y},$$

also $x = \tilde{x}$ und $y = \tilde{y}$ nach der Kürzungsregel in 1.2.3.

Seien umgekehrt die Gleichungen

$$x \cdot a = b \quad \text{und} \quad a \cdot y = b$$

für beliebige $a, b \in G$ lösbar. Dann gibt es zu a ein e mit $e \cdot a = a$. Ist $b \in G$ beliebig, so ist

$$e \cdot b = e \cdot (a \cdot y) = (e \cdot a) \cdot y = a \cdot y = b,$$

also existiert ein neutrales Element. Durch Lösen der Gleichung $x \cdot a = e$ erhält man das inverse Element von a . $\square$

1.2.5. Eine Verknüpfung auf einer endlichen Menge $G = \{a_1, \dots, a_n\}$ kann man im Prinzip dadurch angeben, daß man die Werte aller Produkte $a_i \cdot a_j$ in einem quadratischen Schema ähnlich einer Matrix aufschreibt. Dabei steht $a_i \cdot a_j$ in der i -ten Zeile und der j -ten Spalte der *Verknüpfungstafel* (oder im Fall einer Gruppe der *Gruppentafel*):

*	...	a_j	...
$\vdots$			
a_i		$a_i * a_j$	
$\vdots$			

Ob das Gruppenaxiom G2 erfüllt ist, kann man dann nach obigem Lemma daran erkennen, ob jede Zeile und jede Spalte der Tafel eine Permutation von $a_1, \dots, a_n$ ist.

Daraus folgt sofort, daß es nur eine Möglichkeit gibt, die 2-elementige Menge $G_2 = \{a_1, a_2\}$ zu einer Gruppe zu machen: Ein Element, etwa $a_1 = e$, muß

neutral sein, das andere ist $a_2 = a$, und die Gruppentafel ist

$\cdot$	e	a
e	e	a
a	a	e

Die Kommutativität erkennt man an der Symmetrie der Tafel. Ob das Assoziativgesetz erfüllt ist, kann man der Tafel nicht direkt ansehen, das muß man (leider) für alle möglichen n^3 Tripel nachprüfen.

Für $n = 3$ und $G = \{e, a, b\}$ erhält man ebenfalls nur eine mögliche Gruppentafel, nämlich

$\cdot$	e	a	b
e	e	a	b
a	a	b	e
b	b	e	a

und man stellt fest, daß diese Multiplikation assoziativ ist. Für $n = 4$ und $G = \{e, a, b, c\}$ findet man leicht zwei wesentlich verschiedene Möglichkeiten, nämlich

$\cdot$	e	a	b	c
e	e	a	b	c
a	a	b	c	e
b	b	c	e	a
c	c	e	a	b

und

$\odot$	e	a	b	c
e	e	a	b	c
a	a	e	c	b
b	b	c	e	a
c	c	b	a	e

Wieder ist die Kommutativität offensichtlich, die Assoziativität etwas mühsam zu überprüfen. Um diese beiden verschiedenen Multiplikationen unterscheiden zu können, schreiben wir $G_4^{\odot}$ für G_4 mit der Multiplikation $\odot$.

Es ist klar, daß dieses nur auf Ausprobieren beruhende Verfahren für größere n zu kompliziert und unbefriedigend ist.

1.2.6. Um eine bessere Methode zum Studium von Gruppen zu erhalten, benötigt man geeignete Begriffe, die Beziehungen der Gruppen untereinander regeln.

Definition. Sei G eine Gruppe mit Verknüpfung $\cdot$ und $G' \subset G$ eine nichtleere Teilmenge. G' heißt *Untergruppe*, wenn für $a, b \in G'$ auch $a \cdot b \in G'$ und $a^{-1} \in G'$.

Sind G und H Gruppen mit Verknüpfungen $\cdot$ und $\odot$, so heißt eine Abbildung $\varphi: G \rightarrow H$ *Homomorphismus (von Gruppen)*, wenn

$$\varphi(a \cdot b) = \varphi(a) \odot \varphi(b) \quad \text{für alle } a, b \in G.$$

Ein Homomorphismus heißt *Isomorphismus*, wenn er bijektiv ist.

Zunächst einige unmittelbare Folgerungen aus den Definitionen.

Bemerkung 1. Ist G eine Gruppe und $G' \subset G$ Untergruppe, so ist G' mit der Verknüpfung aus G wieder eine Gruppe.

Man nennt diese Verknüpfung in G' die von G induzierte Verknüpfung.

Beweis. Die induzierte Verknüpfung ist assoziativ, weil sie aus G kommt. Da es ein $a \in G'$ gibt, ist $a^{-1} \in G'$ und $a \cdot a^{-1} = e \in G'$. $\square$

Bemerkung 2. Sei $\varphi: G \rightarrow H$ ein Homomorphismus von Gruppen. Dann gilt:

- a) $\varphi(e) = \mathbf{e}$, wenn $e \in G$ und $\mathbf{e} \in H$ die neutralen Elemente bezeichnen.
- b) $\varphi(a^{-1}) = \varphi(a)^{-1}$ für alle $a \in G$.
- c) Ist φ Isomorphismus, so ist auch die Umkehrabbildung $\varphi^{-1}: H \rightarrow G$ ein Homomorphismus.

Beweis. a) folgt aus der Kürzungsregel in H , da

$$\mathbf{e} \odot \varphi(e) = \varphi(e) = \varphi(e \cdot e) = \varphi(e) \odot \varphi(e),$$

und daraus ergibt sich b), weil

$$\mathbf{e} = \varphi(e) = \varphi(a^{-1} \cdot a) = \varphi(a^{-1}) \odot \varphi(a).$$

Zum Beweis von c) nehmen wir $c, d \in H$. Ist $c = \varphi(a)$ und $d = \varphi(b)$, so ist

$$\varphi(a \cdot b) = \varphi(a) \odot \varphi(b) = c \odot d, \quad \text{also} \quad \varphi^{-1}(c \odot d) = a \cdot b = \varphi^{-1}(c) \cdot \varphi^{-1}(d).$$

$\square$

Im folgenden werden wir die Verknüpfung und die neutralen Elemente in G und H nur noch dann verschieden bezeichnen, wenn das erforderlich ist (etwa in dem Beispiel b) weiter unten).

Beispiele. a) Zunächst betrachten wir die in 1.2.5 konstruierten Beispiele. Für die Mengen gilt

$$G_2 \subset G_3, \quad G_3 \subset G_4 \quad \text{und} \quad G_2 \subset G_4,$$

aber bei keiner der Teilmengen handelt es sich um eine Untergruppe. Dagegen ist

$$G_2 \subset G_4^\circ$$

eine Untergruppe.

Die identische Abbildung $G_4 \rightarrow G_4^\circ$ ist kein Homomorphismus. Beispiele von Homomorphismen $\varphi: G_4 \rightarrow G_4^\circ$ sind gegeben durch

$$\varphi(e) = \varphi(a) = \varphi(b) = \varphi(c) = e,$$

$$\varphi(e) = e, \quad \varphi(a) = a, \quad \varphi(b) = e, \quad \varphi(c) = a,$$

$$\varphi(e) = e, \quad \varphi(a) = b, \quad \varphi(b) = e, \quad \varphi(c) = b,$$

$$\varphi(e) = e, \quad \varphi(a) = c, \quad \varphi(b) = e, \quad \varphi(c) = c.$$

Die einfachen Begründungen seien dem Leser zur Übung empfohlen.

b) Ist $G = \mathbb{R}$ mit der Addition und $H = \mathbb{R}_+^*$ mit der Multiplikation als Verknüpfung, so ist die Exponentialfunktion

$$\exp: \mathbb{R} \rightarrow \mathbb{R}_+^*, \quad x \mapsto e^x,$$

ein Isomorphismus, da $e^{x+y} = e^x \cdot e^y$.

c) Wir betrachten die abelsche Gruppe $\mathbb{Z}$ mit der Addition als Verknüpfung. Für jedes $m \in \mathbb{Z}$ ist die Abbildung

$$\varphi_m: \mathbb{Z} \rightarrow \mathbb{Z}, \quad a \mapsto m \cdot a,$$

ein Homomorphismus, denn $m(a+b) = ma + mb$. Sein Bild

$$m\mathbb{Z} := \{m \cdot a : a \in \mathbb{Z}\} \subset \mathbb{Z}$$

ist eine Untergruppe, weil $ma + mb = m(a+b)$ und $-mb = m(-b)$.

1.2.7. Im Anschluß an das letzte Beispiel kann man nun eine sogenannte *zyklische Gruppe* mit m Elementen konstruieren. Wir nehmen $m > 0$ an und teilen die ganzen Zahlen in m Klassen (d.h. disjunkte Teilmengen) folgendermaßen ein: Zu jedem

$$r \in \{0, 1, \dots, m-1\}$$

betrachten wir die Menge

$$r + m\mathbb{Z} := \{r + m \cdot a : a \in \mathbb{Z}\},$$

die man sich als die um r verschobene Untergruppe $m\mathbb{Z}$ vorstellen kann (vgl. 1.1.8). Für $m = 2$ ist

$0 + 2\mathbb{Z}$ die Menge der geraden und

$1 + 2\mathbb{Z}$ die Menge der ungeraden Zahlen.

Im allgemeinen Fall ist

$$\mathbb{Z} = (0 + m\mathbb{Z}) \cup (1 + m\mathbb{Z}) \cup \dots \cup (m-1 + m\mathbb{Z}),$$

und diese Vereinigung ist disjunkt. Zu welcher Klasse eine Zahl $a \in \mathbb{Z}$ gehört, kann man durch Division mit Rest entscheiden. Ist

$$\frac{a}{m} = k + \frac{r}{m} \quad \text{mit } k \in \mathbb{Z} \text{ und } r \in \{0, \dots, m-1\},$$

so ist $a \in r + m\mathbb{Z}$, denn $a = r + mk$. Daher nennt man die Klassen $r + m\mathbb{Z}$ auch *Restklassen modulo m* : sie bestehen aus all den Zahlen, die bei Division durch m den gleichen Rest hinterlassen. Zwei Zahlen a, a' liegen genau dann in derselben Restklasse, wenn $a - a'$ durch m teilbar ist. Man schreibt dafür auch

$$a \equiv a' \pmod{m}$$

und sagt „ a ist kongruent a' modulo m “. Für $m = 0$ ist die Kongruenz die Gleichheit.

Zu jedem $a \in \mathbb{Z}$ bezeichnen wir mit $\bar{a} = a + m\mathbb{Z}$ seine Restklasse. Die Zahl a heißt *Repräsentant* von $\bar{a}$. In der Menge der Restklassen kann man nun eine Addition definieren durch

$$\bar{a} + \bar{b} := \overline{a + b}.$$

Diese Definition ist nur sinnvoll, wenn sie nicht von der Auswahl der Repräsentanten abhängt. Man sagt dazu, die Addition ist *wohldefiniert*. Das ist aber einfach zu sehen:

Ist $\bar{a} = \bar{a}'$ und $\bar{b} = \bar{b}'$, so ist $a - a' = mk$ und $b - b' = ml$ mit $k, l \in \mathbb{Z}$, also

$$a + b = a' + b' + m(k + l), \quad \text{d.h.} \quad \overline{a + b} = \overline{a' + b'}.$$

Nun können wir das Ergebnis formulieren:

Satz. Ist $m \in \mathbb{Z}$ und $m > 0$, so ist die Menge

$$\mathbb{Z}/m\mathbb{Z} := \{\bar{0}, \bar{1}, \dots, \overline{m-1}\}$$

der Restklassen modulo m mit der oben erklärten Addition eine abelsche Gruppe, und die Abbildung

$$\mathbb{Z} \rightarrow \mathbb{Z}/m\mathbb{Z}, \quad a \mapsto a + m\mathbb{Z},$$

ist ein surjektiver Homomorphismus.

Beweis. Die Assoziativität vererbt sich von $\mathbb{Z}$ nach $\mathbb{Z}/m\mathbb{Z}$:

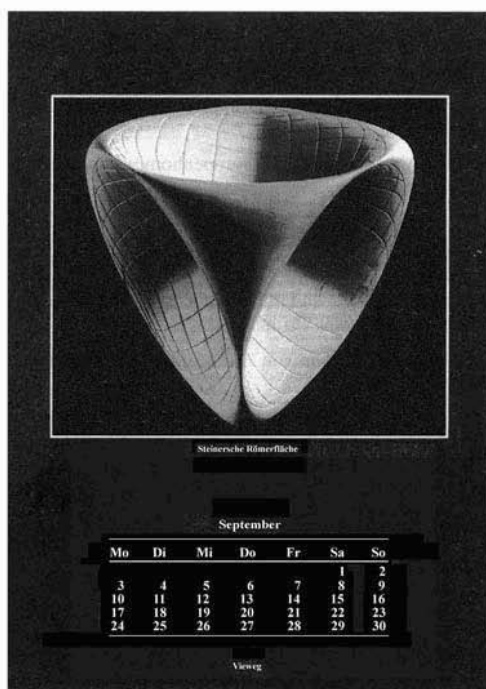
$$\begin{aligned} (\bar{a} + \bar{b}) + \bar{c} &= \overline{a + b} + \bar{c} = \overline{(a + b) + c} = \overline{a + (b + c)} \\ &= \overline{a + b + c} = \bar{a} + \overline{(b + c)}. \end{aligned}$$

Neutrales Element ist $\bar{0}$, denn $\bar{0} + \bar{a} = \overline{0 + a} = \bar{a}$, und Inverses von $\bar{a}$ ist $\overline{-a}$. Auch die Kommutativität wird von $\mathbb{Z}$ vererbt. $\square$

Man nennt $\mathbb{Z}/m\mathbb{Z}$ für $m > 0$ die *zyklische Gruppe der Ordnung m* , für $m = 0$ ist $\mathbb{Z}/0\mathbb{Z} = \mathbb{Z}$, diese Gruppe heißt *unendlich zyklisch*.

Das etwas ungewohnt erscheinende Rechnen mit Restklassen ist im täglichen Leben höchst vertraut: Wer um 10 Uhr weggeht und 3 Stunden unterwegs ist, kommt um 1 Uhr zurück. Denn der Stundenzeiger der Uhr rechnet modulo 12. Die Zeitrechnung insgesamt mit ihren durch Sekunden, Minuten, Stunden, Tage, Monate und Jahre verursachten Kongruenzen ist weit komplizierter als dieser letzte Abschnitt über Gruppen.

Teile der Restklassen modulo 7 findet man auf jedem Blatt eines Monatskalenders. Betrachtet man die Wochentage als Restklassen, so ergibt deren Addition in diesem Monat z.B. Mittwoch + Samstag = Donnerstag. Die ebenfalls abgebildete *Römerfläche* wird in der projektiven Geometrie wieder auftauchen [Fi].



Ein Leser, der mit Kongruenzen immer noch Schwierigkeiten hat, ist in guter Gesellschaft mit GOETHE'S Faust, der zu Mephisto sagt: „Mich dünkt, die Alte spricht im Fieber“, als die Hexe aus ihrem Rechenbuche vorliest:

*Du must verstehn!
Aus Eins mach Zehn,
Und Zwei laß gehn,
Und Drei mach gleich,
So bist du reich.
Verlier die Vier!*

*Aus Fünf und Sechs,
So sagt die Hex',
Mach Sieben und Acht,
So ist's vollbracht:
Und Neun ist Eins,
Und Zehn ist keins.
Das ist das Hexen-Einmaleins.*

Alles klar: die Hexe rechnet schließlich modulo 2. Nur am Anfang holperst noch etwas, da scheint der Reim vor der Rechnung zu stehen.

Aufgaben zu 1.2

1. Sei G eine Gruppe mit $aa = e$ für alle $a \in G$, wobei e das neutrale Element von G bezeichnet. Zeigen Sie, dass G abelsch ist.

2. Bestimmen Sie (bis auf Isomorphie) alle Gruppen mit höchstens vier Elementen. Welche davon sind abelsch?

3. Welche der folgenden Abbildungen sind Gruppenhomomorphismen?

a) $f_1: \mathbb{Z} \rightarrow \mathbb{Z}, z \mapsto 2z$, b) $f_2: \mathbb{Z} \rightarrow \mathbb{Z}, z \mapsto z + 1$,

c) $f_3: \mathbb{Z} \rightarrow \mathbb{Q}^*, z \mapsto z^2 + 1$, d) $f_4: \mathbb{C}^* \rightarrow \mathbb{R}^*, z \mapsto |z|$,

e) $f_5: \mathbb{C} \rightarrow \mathbb{R}^*, z \mapsto |z|$, f) $f_6: \mathbb{Z}/p\mathbb{Z} \rightarrow \mathbb{Z}/p\mathbb{Z}, z \mapsto z^p$.

Dabei ist die Verknüpfung in $\mathbb{Z}$, $\mathbb{C}$ und $\mathbb{Z}/p\mathbb{Z}$ jeweils die Addition, in $\mathbb{Q}^*$, $\mathbb{R}^*$ und $\mathbb{C}^*$ jeweils die Multiplikation und p eine Primzahl.

4. Sei G eine Gruppe und $A \subset G$. Die von A erzeugte Untergruppe $\text{erz}(A)$ ist definiert durch

$$\text{erz}(A) = \{a_1 \cdot \dots \cdot a_n : n \in \mathbb{N}, a_i \in A \text{ oder } a_i^{-1} \in A\}.$$

$\text{erz}(A)$ ist somit die Menge aller endlichen Produkte von Elementen aus A bzw. deren Inversen. Zeigen Sie, dass $\text{erz}(A)$ die „kleinste“ Untergruppe von G ist, die A enthält, d.h.

i) $\text{erz}(A) \subset G$ ist eine Untergruppe.

ii) Ist $U \subset G$ eine Untergruppe mit $A \subset U$, so folgt $\text{erz}(A) \subset U$.

Wie sieht $\text{erz}(A)$ aus für den Fall, dass A einelementig ist?

5. Für eine natürliche Zahl $n \geq 3$ sei $d \in S(\mathbb{R}^2)$ die Drehung um den Winkel $2\pi/n$ und $s \in S(\mathbb{R}^2)$ die Spiegelung an der x -Achse. Die *Diedergruppe* D_n ist definiert durch

$$D_n := \text{erz}(\{s, d\}).$$

a) Wie viele Elemente hat D_n ?

b) Geben Sie eine Gruppentafel von D_3 an.

6. Eine Gruppe G heißt *zyklisch*, falls es ein $g \in G$ gibt mit $G = \text{erz}(\{g\})$.

a) Wie sieht die Gruppentafel einer endlichen zyklischen Gruppe aus?

b)* Zeigen Sie, dass jede zyklische Gruppe entweder isomorph zu $\mathbb{Z}$ oder $\mathbb{Z}/n\mathbb{Z}$ ($n \in \mathbb{N}$ geeignet) ist.

7. Zeigen Sie: Ist G eine abelsche Gruppe und $H \subset G$ eine Untergruppe, so ist durch

$$x \sim y \Leftrightarrow xy^{-1} \in H$$

eine Äquivalenzrelation auf G erklärt. Sei $G/H := G/\sim$ die Menge der Äquivalenzklassen, und die zu $x \in G$ gehörige Äquivalenzklasse sei mit $\bar{x}$ bezeichnet. Sind $x, x', y, y' \in G$ mit $x \sim x'$ und $y \sim y'$, so ist $xy \sim x'y'$. Somit kann man auf G/H durch

$$\bar{x} \cdot \bar{y} := \overline{xy}$$

eine Verknüpfung erklären.

Zeigen Sie, dass G/H auf diese Weise zu einer abelschen Gruppe wird und für $G = \mathbb{Z}$, $H = n\mathbb{Z}$ genau die in 1.2.7 definierten zyklischen Gruppen $\mathbb{Z}/n\mathbb{Z}$ entstehen.

8. Man gebe ein Beispiel einer nicht assoziativen Verknüpfung auf der Menge $G = \{1, 2, 3\}$, so dass für alle $a \in G$ die Translationen τ_a und ${}_a\tau$ aus 1.2.4 surjektiv sind.

1.3 Ringe, Körper und Polynome

1.3.1. Bei Gruppen hat man eine einzige Verknüpfung; ob man sie als Addition oder als Multiplikation bezeichnet, ist nebensächlich. In der linearen Algebra braucht man aber Addition und Multiplikation, also zwei Arten der Verknüpfung, die miteinander in geregelter Beziehung stehen. Damit beschäftigt sich dieser Abschnitt.

Definition. Eine Menge R zusammen mit zwei Verknüpfungen

$$\begin{aligned} +: R \times R &\rightarrow R, & (a, b) &\mapsto a + b, & \text{ und} \\ \cdot: R \times R &\rightarrow R, & (a, b) &\mapsto a \cdot b, \end{aligned}$$

heißt *Ring*, wenn folgendes gilt:

R1 R zusammen mit der Addition $+$ ist eine abelsche Gruppe.

R2 Die Multiplikation $\cdot$ ist assoziativ.

R3 Es gelten die *Distributivgesetze*, d.h. für alle $a, b, c \in R$ gilt

$$a \cdot (b + c) = a \cdot b + a \cdot c \quad \text{und} \quad (a + b) \cdot c = a \cdot c + b \cdot c.$$

Ein Ring heißt *kommutativ*, wenn $a \cdot b = b \cdot a$ für alle $a, b \in R$. Ein Element $1 \in R$ heißt *Einselement*, wenn $1 \cdot a = a \cdot 1 = a$ für alle $a \in R$.

Wie üblich soll dabei zur Einsparung von Klammern die Multiplikation stärker binden als die Addition.

Bemerkung. Ist R ein Ring und $0 \in R$ das Nullelement, so gilt für alle $a \in R$

$$0 \cdot a = a \cdot 0 = 0.$$

Beweis. $0 \cdot a = (0 + 0) \cdot a = 0 \cdot a + 0 \cdot a.$ □

Beispiele. a) Die Mengen $\mathbb{Z}$ der ganzen Zahlen, $\mathbb{Q}$ der rationalen Zahlen und $\mathbb{R}$ der reellen Zahlen sind zusammen mit der üblichen Addition und Multiplikation kommutative Ringe.

b) Ist $I \subset \mathbb{R}$ ein Intervall und

$$R := \{f: I \rightarrow \mathbb{R}\}$$

die Menge der reellwertigen Funktionen, so sind durch

$$(f + g)(x) := f(x) + g(x) \quad \text{und} \quad (f \cdot g)(x) := f(x) \cdot g(x)$$

Verknüpfungen erklärt, und R wird damit zu einem kommutativen Ring. Das folgt ganz leicht aus den Ringeigenschaften von $\mathbb{R}$.

c) In der Gruppe $\mathbb{Z}/m\mathbb{Z}$ der Restklassen modulo m aus 1.2.7 kann man durch

$$\overline{a} \cdot \overline{b} := \overline{a \cdot b}$$

auch eine Multiplikation erklären. Denn ist wieder $a - a' = mk$ und $b - b' = ml$, so folgt

$$a \cdot b = (a' + mk) \cdot (b' + ml) = a' \cdot b' + m(b'k + a'l + mkl).$$

Also ist die Definition der Multiplikation unabhängig von der Auswahl der Repräsentanten.

Die Regeln R2 und R3 und die Kommutativität der Multiplikation folgen ganz einfach aus den entsprechenden Regeln in $\mathbb{Z}$.

Die Multiplikationstafeln wollen wir für $m = 2, 3, 4$ explizit aufschreiben. Dabei lassen wir zur Vereinfachung bei den Restklassen die Querstriche weg.

·	0	1
0	0	0
1	0	1

·	0	1	2
0	0	0	0
1	0	1	2
2	0	2	1

·	0	1	2	3
0	0	0	0	0
1	0	1	2	3
2	0	2	0	2
3	0	3	2	1

Die Multiplikation mit 0 ist uninteressant, wir betrachten also in diesen drei Fällen die Menge $R \setminus \{0\}$. Für $m = 2$ und $m = 3$ ist sie zusammen mit der Multiplikation wieder eine Gruppe, für $m = 4$ nicht, denn hier ist

$$1 \cdot 2 = 3 \cdot 2 \quad \text{und} \quad 2 \cdot 2 = 0.$$

Also ist die Kürzungsregel verletzt, und das Produkt $2 \cdot 2$ liegt nicht in $R \setminus \{0\}$.

1.3.2. Das vorangehende Beispiel motiviert die

Definition. Ein Ring R heißt *nullteilerfrei*, wenn für alle $a, b \in R$ aus $a \cdot b = 0$ stets $a = 0$ oder $b = 0$ folgt.

Bemerkung. Der Restklassenring $\mathbb{Z}/m\mathbb{Z}$ ist genau dann nullteilerfrei, wenn m eine Primzahl ist.

Beweis. Ist m keine Primzahl, also $m = k \cdot l$ mit $1 < k, l < m$, so ist

$$\bar{k}, \bar{l} \neq \bar{0}, \quad \text{aber} \quad \bar{0} = \bar{m} = \bar{k} \cdot \bar{l}.$$

Ist umgekehrt m Primzahl und $\bar{k} \cdot \bar{l} = \bar{0}$, so ist

$$k \cdot l = r \cdot m$$

für ein $r \in \mathbb{Z}$. Also hat entweder k oder l einen Primfaktor m , d.h. $\bar{k} = \bar{0}$ oder $\bar{l} = \bar{0}$. $\square$

Als Vorsorge für später noch eine

Definition. Ist R ein Ring und $R' \subset R$ eine Teilmenge, so heißt R' *Unterring*, wenn R' bezüglich der Addition Untergruppe ist (also entsprechend 1.2.6 für $a, b \in R'$ auch $a + b \in R'$ und $-a \in R'$), und bezüglich der Multiplikation für $a, b \in R'$ auch $a \cdot b \in R'$.

Sind R und S Ringe mit Verknüpfungen $+$, $\cdot$ und $\oplus$, $\odot$, so heißt eine Abbildung $\varphi: R \rightarrow S$ ein *Homomorphismus (von Ringen)*, wenn für alle $a, b \in R$ gilt:

$$\varphi(a + b) = \varphi(a) \oplus \varphi(b) \quad \text{und} \quad \varphi(a \cdot b) = \varphi(a) \odot \varphi(b).$$

Zum Beispiel ist $m\mathbb{Z} \subset \mathbb{Z}$ ein Unterring und $\mathbb{Z} \rightarrow \mathbb{Z}/m\mathbb{Z}$, $a \mapsto a + m\mathbb{Z}$, ein Homomorphismus.

1.3.3. In einem nullteilerfreien Ring R ist für $a, b \in R \setminus \{0\}$ auch das Produkt $a \cdot b \in R \setminus \{0\}$, also induziert die Multiplikation von R eine assoziative Multiplikation in $R \setminus \{0\}$. Das Gruppenaxiom G2 braucht jedoch keineswegs erfüllt zu sein: Im Ring $\mathbb{Z}$ gibt es außer für 1 und -1 kein multiplikatives Inverses, im Ring $2\mathbb{Z}$ der geraden Zahlen nicht einmal ein Einselement. Ist $R \setminus \{0\}$ mit der Multiplikation eine Gruppe und darüber hinaus der Ring kommutativ, so nennt man ihn *Körper*. Das wollen wir noch einmal direkter aufschreiben:

Definition. Eine Menge K zusammen mit zwei Verknüpfungen

$$\begin{aligned} +: K \times K &\rightarrow K, & (a, b) &\mapsto a + b, & \text{und} \\ \cdot: K \times K &\rightarrow K, & (a, b) &\mapsto a \cdot b, \end{aligned}$$

heißt *Körper*, wenn folgendes gilt:

K1 K zusammen mit der Addition $+$ ist eine abelsche Gruppe.

(Ihr neutrales Element wird mit 0, das zu $a \in K$ inverse Element mit $-a$ bezeichnet.)

K2 Bezeichnet $K^* := K \setminus \{0\}$, so gilt für $a, b \in K^*$ auch $a \cdot b \in K^*$, und K^* zusammen mit der so erhaltenen Multiplikation ist eine abelsche Gruppe.

(Ihr neutrales Element wird mit 1, das zu $a \in K^*$ inverse Element mit a^{-1} oder $1/a$ bezeichnet. Man schreibt $b/a = a^{-1}b = ba^{-1}$.)

K3 Es gelten die *Distributivgesetze*, d.h. für $a, b, c \in K$ ist

$$a \cdot (b + c) = a \cdot b + a \cdot c \quad \text{und} \quad (a + b) \cdot c = a \cdot c + b \cdot c.$$

Bemerkung. In einem Körper K gelten die folgenden weiteren Rechenregeln (dabei sind $a, b, x, \tilde{x} \in K$ beliebig):

a) $1 \neq 0$ (also hat ein Körper mindestens zwei Elemente).

$$b) 0 \cdot a = a \cdot 0 = 0.$$

$$c) a \cdot b = 0 \Rightarrow a = 0 \text{ oder } b = 0.$$

$$d) a(-b) = -(ab) \text{ und } (-a)(-b) = ab.$$

$$e) x \cdot a = \tilde{x} \cdot a \text{ und } a \neq 0 \Rightarrow x = \tilde{x}.$$

Beweis. a) ist ganz klar, denn $1 \in K^*$, aber $0 \notin K^*$. b) sieht man wie in 1.3.1. Die Nullteilerfreiheit c) ist in K2 enthalten. d) folgt aus

$$ab + (-a)b = (a + (-a))b = 0 \cdot b = 0 \quad \text{und}$$

$$(-a)(-b) = -((-a)b) = -(-ab) = ab \quad \text{nach 1.2.3, Bem. c).}$$

Die Kürzungsregel e) gilt in K^* , also im Fall $x, \tilde{x} \in K^*$. Ist $x = 0$, so muß auch $\tilde{x} = 0$, also $x = \tilde{x}$ sein. $\square$

1.3.4. Beispiele für Körper. a) Die rationalen Zahlen $\mathbb{Q}$ und die reellen Zahlen $\mathbb{R}$ sind Körper. Das lernt man in der Analysis (vgl. etwa [Fo1], §2).

b) Zur Konstruktion der komplexen Zahlen $\mathbb{C}$ führt man in der reellen Ebene $\mathbb{R} \times \mathbb{R}$ eine Addition und Multiplikation ein. Die naheliegende Frage, warum das im $\mathbb{R}^n$ mit $n > 2$ nicht mehr so geht, wird in [Z] behandelt.

Durch einfaches Nachprüfen der Körperaxiome beweist man:

$$\mathbb{R} \times \mathbb{R} = \{(a, b) : a, b \in \mathbb{R}\}$$

zusammen mit der durch

$$(a, b) + (a', b') := (a + a', b + b')$$

definierten Addition und der durch

$$(a, b) \cdot (a', b') := (aa' - bb', ab' + a'b)$$

definierten Multiplikation ist ein Körper mit $(0,0)$ als neutralem Element der Addition, $(-a, -b)$ als Negativem von (a, b) , $(1,0)$ als neutralem Element der Multiplikation und

$$(a, b)^{-1} := \left(\frac{a}{a^2 + b^2}, \frac{-b}{a^2 + b^2} \right)$$

als multiplikativem Inversen.

Wir nennen $\mathbb{R} \times \mathbb{R}$ mit diesen Verknüpfungen den Körper der komplexen Zahlen und bezeichnen ihn mit $\mathbb{C}$.

Die Abbildung

$$\mathbb{R} \rightarrow \mathbb{R} \times \mathbb{R} = \mathbb{C}, \quad a \mapsto (a, 0),$$

ist injektiv. Da

$$\begin{aligned}(a, 0) + (a', 0) &= (a + a', 0) \quad \text{und} \\ (a, 0) \cdot (a', 0) &= (a \cdot a', 0)\end{aligned}$$

gilt, braucht man zwischen $\mathbb{R}$ und

$$\mathbb{R} \times \{0\} = \{(a, b) \in \mathbb{C} : b = 0\}$$

auch hinsichtlich Addition und Multiplikation nicht zu unterscheiden. Man kann also $\mathbb{R}$ mit $\mathbb{R} \times \{0\}$, d.h. a mit $(a, 0)$ „identifizieren“ und $\mathbb{R}$ als Teilmenge von $\mathbb{C}$ betrachten.

Dies wird noch einleuchtender durch folgende übliche Konventionen. Man definiert $i := (0, 1)$ als *imaginäre Einheit*. Dann ist $i^2 = -1$, und für jedes $(a, b) \in \mathbb{C}$ gilt

$$(a, b) = (a, 0) + (0, b) = a + bi.$$

Für $\lambda = (a, b) = a + bi \in \mathbb{C}$ nennt man $\operatorname{re} \lambda := a \in \mathbb{R}$ den *Realteil* und $\operatorname{im} \lambda := b \in \mathbb{R}$ den *Imaginärteil*,

$$\bar{\lambda} := a - bi \in \mathbb{C}$$

heißt die zu λ *konjugiert komplexe Zahl*.

Für die komplexe Konjugation gelten folgende, ganz einfach nachzuweisende Regeln: Für alle $\lambda, \mu \in \mathbb{C}$ ist

$$\begin{aligned}\overline{\lambda + \mu} &= \bar{\lambda} + \bar{\mu}, \\ \overline{\lambda \cdot \mu} &= \bar{\lambda} \cdot \bar{\mu}, \\ \lambda \in \mathbb{R} &\Leftrightarrow \lambda = \bar{\lambda}.\end{aligned}$$

Da für $\lambda = a + bi \in \mathbb{C}$

$$\lambda \cdot \bar{\lambda} = (a + bi) \cdot (a - bi) = a^2 + b^2 \in \mathbb{R}_+,$$

kann man den *Absolutbetrag*

$$|\lambda| := \sqrt{\lambda \cdot \bar{\lambda}} = \sqrt{a^2 + b^2}$$

definieren. Wie man leicht nachrechnet, ist für alle $\lambda, \mu \in \mathbb{C}$

$$|\lambda + \mu| \leq |\lambda| + |\mu| \quad \text{Dreiecksungleichung} \quad \text{und} \quad |\lambda \cdot \mu| = |\lambda| \cdot |\mu|.$$

Vorsicht! Die in $\mathbb{R}$ vorhandene $\leq$ -Relation läßt sich nicht in sinnvoller Weise auf $\mathbb{C}$ fortsetzen. Für komplexe Zahlen kann man daher i.a. nur die Absolutbeträge vergleichen, d.h. für $\lambda, \mu \in \mathbb{C}$ ist

$$|\lambda| \leq |\mu| \quad \text{oder} \quad |\mu| \leq |\lambda|.$$

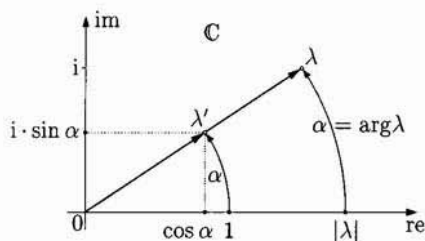


Bild 1.4

Wir wollen noch eine geometrische Beschreibung von Addition und Multiplikation komplexer Zahlen geben. Die Addition entspricht der Addition von Vektoren im $\mathbb{R}^2$ (Bild 1.5, vgl. auch Bild 0.1). Ist λ eine von Null verschiedene komplexe Zahl und $\lambda' = \frac{1}{|\lambda|} \cdot \lambda$, so ist $|\lambda'| = 1$. Es gibt also ein eindeutig bestimmtes $\alpha \in [0, 2\pi[$, so daß

$$\lambda' = \cos \alpha + i \cdot \sin \alpha = e^{i\alpha},$$

wie man in der Analysis lernt.

Man nennt $\arg \lambda := \alpha$ das *Argument* von λ , und es ist

$$\lambda = |\lambda| \cdot e^{i \arg \lambda}.$$

Ist $\mu = |\mu| \cdot e^{i \arg \mu}$ eine weitere von Null verschiedene komplexe Zahl, so ist

$$\lambda \cdot \mu = |\lambda| \cdot |\mu| \cdot e^{i \arg \lambda} \cdot e^{i \arg \mu} = |\lambda| \cdot |\mu| \cdot e^{i(\arg \lambda + \arg \mu)}.$$

Bei der Multiplikation komplexer Zahlen werden also die Absolutbeträge multipliziert und die Argumente addiert (Bild 1.5).

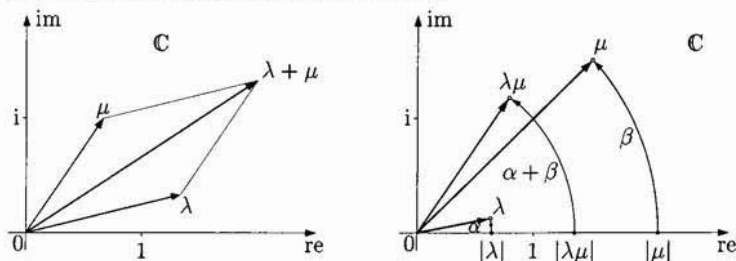


Bild 1.5

c) Wie wir gesehen haben, gibt es in jedem Körper zwei verschiedene Elemente 0 und 1. Daraus kann man schon einen Körper machen, indem man in $K = \{0, 1\}$ Addition und Multiplikation einführt durch die Tafeln

+	0	1
0	0	1
1	1	0

·	0	1
0	0	0
1	0	1

Offensichtlich ist das auch die einzige Möglichkeit, als Ring war dieser Körper schon in 1.3.1 in der Form $\mathbb{Z}/2\mathbb{Z}$ aufgetreten. Diese Verknüpfungen kann man elektronisch leicht realisieren, daher ist dieser Körper der elementare Baustein aller Computer.

d) In 1.3.2 hatten wir für jedes $m \in \mathbb{N} \setminus \{0\}$ den Restklassenring $\mathbb{Z}/m\mathbb{Z}$ eingeführt und bewiesen, daß er genau dann nullteilerfrei ist, wenn m eine Primzahl ist. Dies ist also eine notwendige Bedingung dafür, ein Körper zu sein. Daß es auch hinreichend ist, folgt aus der etwas allgemeineren

Bemerkung. Ein nullteilerfreier, kommutativer Ring K mit endlich vielen Elementen und Eins ist ein Körper.

Beweis. Nach 1.2.4 genügt es, für jedes $a \in K^*$ zu zeigen, daß die Multiplikation

$$K^* \rightarrow K^*, \quad x \mapsto a \cdot x,$$

eine surjektive Abbildung ist. Wenn K und damit auch K^* endlich ist, genügt dafür die Injektivität (vgl. 1.1.4). Die ist aber klar, denn für $x \neq \tilde{x}$ und $a \cdot x = a \cdot \tilde{x}$ würde

$$a(x - \tilde{x}) = 0 \quad \text{und} \quad a \neq 0, \quad x - \tilde{x} \neq 0$$

gelten. □

Im Ring $\mathbb{Z}$ gilt für jedes $n \in \mathbb{N}$ mit $n \geq 1$

$$n \cdot 1 := \underbrace{1 + \dots + 1}_{n\text{-mal}} = n \neq 0.$$

In $\mathbb{Z}/m\mathbb{Z}$ mit $m \geq 1$ ist die Restklasse $\bar{1}$ das Einselement, es gilt

$$m \cdot \bar{1} := \underbrace{\bar{1} + \dots + \bar{1}}_{m\text{-mal}} = \bar{1} + \dots + \bar{1} = \bar{m} = \bar{0}.$$

Das zeigt einen grundlegenden Unterschied zwischen den beiden Ringen $\mathbb{Z}$ und $\mathbb{Z}/m\mathbb{Z}$ und motiviert die

Definition. Ist R ein Ring mit Einselement 1, so ist seine *Charakteristik* erklärt durch

$$\text{char}(R) := \begin{cases} 0, & \text{falls } n \cdot 1 \neq 0 \text{ für alle } n \geq 1, \\ \min \{n \in \mathbb{N} \setminus \{0\} : n \cdot 1 = 0\} & \text{sonst.} \end{cases}$$

Lemma. Ist K ein Körper, so ist $\text{char}(K)$ entweder Null oder eine Primzahl.

Beweis. Angenommen, $\text{char}(K) = m = k \cdot \ell \neq 0$ mit $1 < k, \ell < m$. Aus

$$0 = m \cdot 1 = (k \cdot \ell) \cdot 1 = (k \cdot 1)(\ell \cdot 1)$$

folgt wegen der Nullteilerfreiheit $k \cdot 1 = 0$ oder $\ell \cdot 1 = 0$ im Widerspruch zur Minimalität von m . $\square$

Ist p eine Primzahl, so nennt man $\mathbb{Z}/p\mathbb{Z}$ den *Primkörper der Charakteristik p* . In der Algebra lernt man, daß es zu jedem endlichen Körper K eine Primzahl p gibt derart, daß $\mathbb{Z}/p\mathbb{Z}$ Unterkörper von K ist, und die Anzahl der Elemente von K eine Potenz von p ist.

1.3.5. Spätestens bei der Suche nach Eigenwerten in Kapitel 4 wird es sich nicht mehr vermeiden lassen, Polynome beliebigen Grades zu Hilfe zu nehmen. Weil es von der Systematik paßt, wollen wir die dann benötigten Tatsachen schon hier zusammenstellen.

Wir nehmen einen Körper K und eine *Unbestimmte* t . Eine Unbestimmte soll dabei einfach ein Buchstabe sein, für den man alles einsetzen darf, was sinnvoll ist (das kann man präziser formulieren, aber damit wollen wir uns hier nicht aufhalten, vgl. [F-S]). Ein *Polynom* mit *Koeffizienten* in K (oder *Polynom über K*) ist dann ein formaler Ausdruck der Gestalt

$$f(t) = a_0 + a_1 t + \dots + a_n t^n,$$

wobei $a_0, \dots, a_n \in K$. Meist schreibt man statt $f(t)$ nur f . Mit $K[t]$ bezeichnen wir die Menge all solcher Polynome. Sind alle Koeffizienten $a_v = 0$, so spricht man vom *Nullpolynom* und schreibt $f = 0$. Der *Grad* von f ist erklärt als

$$\deg f := \begin{cases} -\infty, & \text{falls } f = 0, \\ \max\{v \in \mathbb{N} : a_v \neq 0\}, & \text{sonst.} \end{cases}$$

Schließlich heißt f *normiert*, wenn $a_n = 1$.

Das nächstliegende, was man für die Unbestimmte t einsetzen kann, sind Elemente aus K . Ist $\lambda \in K$, so ist auch

$$f(\lambda) := a_0 + a_1 \lambda + \dots + a_n \lambda^n \in K,$$

aus dem Polynom f erhält man also eine Abbildung

$$\tilde{f}: K \rightarrow K, \quad \lambda \mapsto f(\lambda),$$

insgesamt also (mit der Notation aus 1.1.3) eine Abbildung

$$\tilde{} : K[t] \rightarrow \text{Abb}(K, K), \quad f \mapsto \tilde{f}.$$

Die etwas pedantische Unterscheidung von f und $\tilde{f}$ ist leider nötig, wenn man sich einmal mit endlichen Körpern eingelassen hat.

Beispiel. Ist $K = \{0, 1\}$ der Körper mit zwei Elementen aus 1.3.4 und

$$f = t^2 + t, \quad \text{so ist} \quad f(0) = 0 + 0 = 0 \quad \text{und} \quad f(1) = 1 + 1 = 0,$$

also $\tilde{f}$ die Nullabbildung, obwohl $f \neq 0$, weil $a_1 = a_2 = 1$. Die obige Abbildung $\tilde{}$ ist also nicht injektiv.

1.3.6. Die Menge $K[t]$ hat viel Struktur, insbesondere eine natürliche Addition und Multiplikation. Dazu nehmen wir $f, g \in K[t]$. Ist

$$f = a_0 + a_1 t + \dots + a_n t^n, \quad g = b_0 + b_1 t + \dots + b_m t^m,$$

so können wir zur Definition der Addition $m = n$ annehmen (ist etwa $m < n$, so setze man $b_{m+1} = \dots = b_n = 0$). Dann ist

$$f + g := (a_0 + b_0) + (a_1 + b_1)t + \dots + (a_n + b_n)t^n.$$

Die Multiplikation ist dadurch erklärt, daß man formal ausmultipliziert, also

$$f \cdot g := c_0 + c_1 t + \dots + c_{n+m} t^{n+m} \quad \text{mit} \quad c_k = \sum_{i+j=k} a_i b_j.$$

Insbesondere ist

$$\begin{aligned} c_0 &= a_0 b_0, \\ c_1 &= a_0 b_1 + a_1 b_0, \\ c_2 &= a_0 b_2 + a_1 b_1 + a_2 b_0, \\ &\vdots \\ c_{n+m} &= a_n b_m. \end{aligned}$$

Ist $f \cdot g = h$, so nennt man f und g *Teiler* von h .

Bemerkung. Ist K ein Körper, so ist die Menge $K[t]$ der Polynome über K zusammen mit den oben definierten Verknüpfungen ein kommutativer Ring. Weiter gilt

$$\deg(f \cdot g) = \deg f + \deg g$$

für $f, g \in K[t]$. Dabei soll formal $n - \infty = -\infty + m = -\infty + (-\infty) = -\infty$ sein.

Man nennt $K[t]$ den *Polynomring* über K .

Beweis. Der Nachweis der Axiome erfordert nur geduldiges Rechnen. Die Aussage über den Grad folgt aus $a_n b_m \neq 0$, falls $a_n, b_m \neq 0$. $\square$

Es sei angemerkt, daß man analog für einen kommutativen Ring R einen kommutativen Polynomring $R[t]$ erhält. Die Aussage über den Grad des Produktpolynoms gilt nur über einem nullteilerfreien Ring.

1.3.7. Der Mangel eines Ringes gegenüber einem Körper ist der, daß man im allgemeinen nicht dividieren kann. Bei ganzen Zahlen hat man als Ersatz eine *Division mit Rest* (vgl. 1.2.7), die ganz analog auch für Polynome erklärt werden kann.

Satz. Sind $f, g \in K[t]$, und ist $g \neq 0$, so gibt es dazu eindeutig bestimmte Polynome $q, r \in K[t]$ derart, daß

$$f = q \cdot g + r \quad \text{und} \quad \deg r < \deg g.$$

Man kann die Beziehung auch in der nicht ohne weiteres erlaubten, aber sehr suggestiven Form

$$\frac{f}{g} = q + \frac{r}{g}$$

schreiben. Der Buchstabe q steht für „Quotient“, r für „Rest“.

Beweis. Wir zeigen zunächst die Eindeutigkeit. Seien $q, r, q', r' \in K[t]$ mit

$$f = q \cdot g + r = q' \cdot g + r', \quad \text{wobei} \quad \deg r, \deg r' < \deg g.$$

Durch Subtraktion folgt

$$0 = (q - q') \cdot g + (r - r'), \quad \text{also} \quad (q - q') \cdot g = r' - r. \quad (*)$$

Da zwei gleiche Polynome gleichen Grad haben, folgt aus $q \neq q'$ wegen $g \neq 0$

$$\deg(r' - r) = \deg(q - q') + \deg g \geq \deg g,$$

was nicht sein kann. Also ist $q = q'$, und aus $(*)$ folgt auch $r = r'$.

Zum Beweis der Existenz von q und r geben wir ein Verfahren an, mit dem man diese Polynome schrittweise berechnen kann. Dazu schreibt man die Polynome am besten nach fallenden Potenzen, also

$$f = a_n t^n + \dots + a_1 t + a_0, \quad g = b_m t^m + \dots + b_1 t + b_0,$$

wobei $a_n, b_m \neq 0$. Ist $n < m$, so kann man $q = 0$ und $r = f$ wählen, denn

$$f = 0 \cdot g + f \quad \text{und} \quad \deg f < \deg g.$$

Im Fall $n \geq m$ teilt man zunächst die höchsten Terme von f und g , das ergibt

$$q_1 := \frac{a_n}{b_m} t^{n-m}$$

als höchsten Term von q . Im nächsten Schritt betrachtet man

$$f_1 := f - q_1 \cdot g.$$

Nach Definition ist $\deg f_1 < \deg f$. Ist $\deg f_1 < m$, so kann man $q = q_1$ und $r = f_1$ setzen. Andernfalls wiederholt man den ersten Schritt mit f_1 statt f , d.h. man erhält den nächsten Term q_2 von q und

$$f_2 := f_1 - q_2 \cdot g \quad \text{mit} \quad \deg f_2 < \deg f_1.$$

Da die Grade der f_i bei jedem Schritt um mindestens eins abnehmen, erhält man schließlich ein $k \leq n - m + 1$, so daß für

$$f_k := f_{k-1} - q_k \cdot g \quad \text{erstmal} \quad \deg f_k < \deg g,$$

und damit bricht das Verfahren ab: Setzt man die Gleichungen ineinander ein, so ergibt sich

$$f = q_1 g + f_1 = (q_1 + q_2)g + f_2 = \dots = (q_1 + \dots + q_k)g + f_k,$$

also ist eine Lösung unseres Problems gegeben durch

$$q := q_1 + \dots + q_k \quad \text{und} \quad r := f_k. \quad \square$$

Da bei der Konstruktion von q immer nur durch b_m dividiert wird, kann man im Fall $b_m = 1$ den Körper K durch einen Ring ersetzen.

Beispiel. Sei $K = \mathbb{R}$, $f = 3t^3 + 2t + 1$, $g = t^2 - 4t$.

Die Rechnung verläuft nach folgendem Schema:

$$\begin{array}{r}
 (3t^3 \quad + 2t + 1) : (t^2 - 4t) = 3t + 12 + (50t + 1) : (t^2 - 4t) \\
 \underline{-3t^3 \quad + 12t^2} \\
 12t^2 \quad + 2t + 1 \\
 \underline{-12t^2 \quad + 48t} \\
 50t + 1
 \end{array}$$

Es ist also $q = 3t + 12$ und $r = 50t + 1$.

1.3.8. Nach diesen formalen Vorbereitungen kommen wir nun zum eigentlichen Thema, nämlich der Frage nach der Existenz von *Nullstellen* (d.h. $\lambda \in K$ mit $f(\lambda) = 0$) bei Polynomen. Darauf kann man nämlich viele andere Fragen zurückführen, etwa in Kapitel 4 die Frage nach der Existenz von Eigenwerten. Neben der reinen Existenzfrage ist es für die Praxis wichtig, Verfahren für die näherungsweise Berechnung von Nullstellen zu haben. Das lernt man in der numerischen Mathematik (vgl. etwa [O]).

Beispiele. a) Ist $K = \mathbb{R}$ und $f = t^2 + 1$, so ist $f(\lambda) \geq 1$ für alle $\lambda \in \mathbb{R}$. Also gibt es keine Nullstelle.

b) Ist $K = \{a_0, a_1, \dots, a_n\}$ ein endlicher Körper und

$$f = (t - a_0) \cdot \dots \cdot (t - a_n) + 1,$$

so ist $f(\lambda) = 1$ für alle $\lambda \in K$, also hat f keine Nullstelle.

Zum Glück sind diese beiden Beispiele von ausgewählter Bosheit; im allgemeinen hat man doch etwas mehr Chancen, eine Nullstelle zu finden. Hat man eine gefunden, so genügt es zur Suche nach weiteren, ein Polynom von kleinerem Grad zu betrachten:

Lemma. Ist $\lambda \in K$ eine Nullstelle von $f \in K[t]$, so gibt es ein eindeutig bestimmtes $g \in K[t]$ mit folgenden Eigenschaften:

1) $f = (t - \lambda) \cdot g.$

2) $\deg g = (\deg f) - 1.$

Beweis. Wir dividieren f durch $(t - \lambda)$ mit Rest; es gibt also eindeutig bestimmte $g, r \in K[t]$ mit

$$f = (t - \lambda)g + r \quad \text{und} \quad \deg r < \deg(t - \lambda) = 1.$$

Also ist $r = a_0$ mit $a_0 \in K$. Aus $f(\lambda) = 0$ folgt durch Einsetzen von λ

$$0 = (\lambda - \lambda) \cdot g(\lambda) + r = 0 + a_0,$$

also ist $a_0 = r = 0$, und 1) ist bewiesen. Wegen

$$\deg f = \deg(t - \lambda) + \deg g = 1 + \deg g$$

folgt 2). □

Korollar 1. Sei K ein beliebiger Körper, $f \in K[t]$ ein Polynom und k die Anzahl der Nullstellen von f . Ist f vom Nullpolynom verschieden, so gilt

$$k \leq \deg f.$$

Beweis. Wir führen Induktion über den Grad von f . Für $\deg f = 0$ ist $f = a_0 \neq 0$ ein konstantes Polynom. Dieses hat gar keine Nullstelle, also ist unsere Behauptung richtig.

Sei $\deg f = n \geq 1$, und sei die Aussage schon für alle Polynome $g \in K[t]$ mit $\deg g \leq n - 1$ bewiesen. Wenn f keine Nullstelle hat, ist die Behauptung richtig. Ist $\lambda \in K$ eine Nullstelle, so gibt es nach dem Lemma ein $g \in K[t]$ mit

$$f = (t - \lambda) \cdot g \quad \text{und} \quad \deg g = n - 1.$$

Alle von λ verschiedenen Nullstellen von f müssen auch Nullstellen von g sein. Ist l die Anzahl der Nullstellen von g , so ist nach Induktionsannahme

$$l \leq n - 1, \quad \text{also} \quad k \leq l + 1 \leq n.$$

□

Korollar 2. Ist K unendlich, so ist die Abbildung

$$\sim : K[t] \rightarrow \text{Abb}(K, K), \quad f \mapsto \tilde{f},$$

injektiv.

Beweis. Seien $f_1, f_2 \in K[t]$ und $g := f_2 - f_1$. Ist $\tilde{f}_1 = \tilde{f}_2$, so folgt $\tilde{g} = 0$, d.h. $g(\lambda) = 0$ für alle $\lambda \in K$. Also hat g unendlich viele Nullstellen, und aus Korollar 1 folgt $g = 0$, somit ist $f_1 = f_2$. $\square$

Ist λ Nullstelle von f , also $f = (t - \lambda) \cdot g$, so kann λ auch Nullstelle von g sein. Man spricht dann von einer *mehrfachen* Nullstelle.

Definition. Ist $f \in K[t]$ vom Nullpolynom verschieden und $\lambda \in K$, so heißt

$$\mu(f; \lambda) := \max\{r \in \mathbb{N} : f = (t - \lambda)^r \cdot g \text{ mit } g \in K[t]\}$$

die *Vielfachheit der Nullstelle* λ von f .

Nach dem Lemma gilt $\mu(f; \lambda) = 0 \Leftrightarrow f(\lambda) \neq 0$. Ist

$$f = (t - \lambda)^r \cdot g \quad \text{mit} \quad r = \mu(f; \lambda),$$

so folgt $g(\lambda) \neq 0$. Die Vielfachheit der Nullstelle λ gibt also an, wie oft der Linearfaktor $(t - \lambda)$ in f enthalten ist.

Ist $K = \mathbb{R}$ oder $\mathbb{C}$, so kann man die Vielfachheit der Nullstelle mit den r -ten Ableitungen $f^{(r)}$ von f in Beziehung bringen. Es gilt

$$\mu(f; \lambda) = \max\{r \in \mathbb{N} : f(\lambda) = f'(\lambda) = \dots = f^{(r-1)}(\lambda) = 0\},$$

wie man leicht nachrechnet.

1.3.9. Sind $\lambda_1, \dots, \lambda_k \in K$ die verschiedenen Nullstellen eines Polynoms $f \in K[t]$, und ist $r_i = \mu(f; \lambda_i)$, so ist

$$f = (t - \lambda_1)^{r_1} \cdot \dots \cdot (t - \lambda_k)^{r_k} \cdot g,$$

wobei g ein Polynom vom Grad $n - (r_1 + \dots + r_k)$ ohne Nullstellen ist. Der schlimmste Fall ist $g = f$, der beste $\deg g = 0$, d.h. f zerfällt in *Linearfaktoren*. Die wichtigste Existenzaussage für Nullstellen von Polynomen macht der sogenannte

Fundamentalsatz der Algebra. Jedes Polynom $f \in \mathbb{C}[t]$ mit $\deg f > 0$ hat mindestens eine Nullstelle.

Dieser Satz wurde von C.F. GAUSS erstmals 1799 bewiesen. Es gibt dafür sehr viele Beweise, die aber alle Hilfsmittel aus der Analysis benutzen, denn $\mathbb{C}$ entsteht aus $\mathbb{R}$, und die reellen Zahlen sind ein Produkt der Analysis. Der wohl kürzeste Beweis verwendet Hilfsmittel aus der Theorie der holomorphen Funktionen (vgl. etwa [F-L]).

Hat $f \in \mathbb{C}[t]$ eine Nullstelle λ , so kann man sie herausdividieren, also

$$f = (t - \lambda) \cdot g$$

schreiben. Ist $\deg g > 0$, so hat auch g eine komplexe Nullstelle, und indem man das Verfahren so lange wiederholt, bis das verbleibende Polynom den Grad Null hat, ergibt sich das

Korollar. Jedes Polynom $f \in \mathbb{C}[t]$ zerfällt in Linearfaktoren, d.h. es gibt a und $\lambda_1, \dots, \lambda_n \in \mathbb{C}$ mit $n = \deg f$, so daß

$$f = a(t - \lambda_1) \cdot \dots \cdot (t - \lambda_n).$$

1.3.10. Nun wollen wir aus dem Fundamentalsatz der Algebra Aussagen über die Nullstellen reeller Polynome folgern. Dazu betrachten wir $\mathbb{R}$ als Teilmenge von $\mathbb{C}$ (vgl. 1.3.4). Dann ist auch $\mathbb{R}[t]$ Teilmenge von $\mathbb{C}[t]$.

Lemma. Ist $f \in \mathbb{R}[t]$ und $\lambda \in \mathbb{C}$ eine Nullstelle von f , so ist auch die konjugiert komplexe Zahl $\bar{\lambda} \in \mathbb{C}$ eine Nullstelle von f . Es gilt sogar

$$\mu(f; \lambda) = \mu(f; \bar{\lambda}).$$

Die komplexen Nullstellen eines Polynoms mit reellen Koeffizienten liegen also symmetrisch zur reellen Achse.

Beweis. Ist

$$f = a_0 + a_1 t + \dots + a_n t^n,$$

so ist wegen $a_0 = \bar{a}_0, \dots, a_n = \bar{a}_n$ nach den Rechenregeln für die komplexe Konjugation

$$f(\bar{\lambda}) = a_0 + a_1 \bar{\lambda} + \dots + a_n (\bar{\lambda})^n = \bar{a}_0 + \overline{a_1 \lambda} + \dots + \overline{a_n \lambda^n} = \overline{f(\lambda)} = \bar{0} = 0.$$

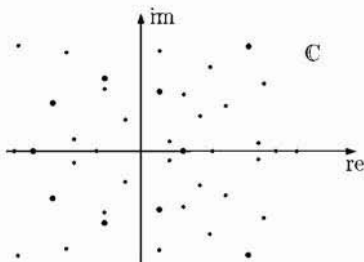


Bild 1.6

Also ist auch $\bar{\lambda}$ Nullstelle von f . Um die Vielfachheiten von λ und $\bar{\lambda}$ zu vergleichen, genügt es für jedes $k \in \mathbb{N}$

$$\mu(f; \lambda) \geq k \Rightarrow \mu(f; \bar{\lambda}) \geq k$$

zu beweisen. Für $\lambda = \bar{\lambda}$ ist die Aussage trivial. Für $\lambda \neq \bar{\lambda}$ verwenden wir den folgenden

Hilfssatz. Sei $f \in \mathbb{R}[t]$ und $\lambda \in \mathbb{C}$ eine nicht reelle Nullstelle von f , sowie $g := (t - \lambda)(t - \bar{\lambda}) \in \mathbb{C}[t]$. Dann gilt:

1) $g \in \mathbb{R}[t]$.

2) Es gibt ein $q \in \mathbb{R}[t]$ mit $f = g \cdot q$.

Nun genügt es, durch Induktion über k aus $\mu(f; \lambda) \geq k$ die Existenz eines Polynoms $f_k \in \mathbb{R}[t]$ mit

$$f = g^k \cdot f_k$$

zu folgern. Dabei ist g wie im Hilfssatz erklärt.

Für $k = 0$ ist nichts zu beweisen. Sei also die Aussage für $k \geq 0$ bewiesen und $\mu(f; \lambda) \geq k + 1$. Dann ist $f = g^k \cdot f_k$, und es muß $f_k(\lambda) = 0$ sein; aus dem Hilfssatz folgt die Existenz eines $f_{k+1} \in \mathbb{R}[t]$ mit

$$f_k = g \cdot f_{k+1}, \quad \text{also ist} \quad f = g^{k+1} \cdot f_{k+1}.$$

Es bleibt der Hilfssatz zu beweisen. Dazu setzen wir

$$\lambda = \alpha + i\beta \quad \text{mit} \quad \alpha, \beta \in \mathbb{R}.$$

Dann ist

$$g = (t - \lambda)(t - \bar{\lambda}) = (t - \alpha - i\beta)(t - \alpha + i\beta) = t^2 - 2\alpha t + \alpha^2 + \beta^2 \in \mathbb{R}[t].$$

Durch Division mit Rest in $\mathbb{R}[t]$ erhält man $q, r \in \mathbb{R}[t]$ mit

$$f = g \cdot q + r \quad \text{und} \quad \deg r \leq (\deg g) - 1 = 1.$$

Diese Gleichung gilt selbstverständlich auch in $\mathbb{C}[t]$. Durch Einsetzen von λ und $\bar{\lambda}$ folgt

$$r(\lambda) = r(\bar{\lambda}) = 0.$$

Nach Korollar 1 aus 1.3.8 folgt $r = 0$ wegen $\lambda \neq \bar{\lambda}$. Damit ist der Hilfssatz und auch das Lemma bewiesen. $\square$

Nun wenden wir auf ein Polynom $f \in \mathbb{R}[t]$ den Fundamentalsatz der Algebra an. Danach gibt es $a, \lambda_1, \dots, \lambda_n \in \mathbb{C}$, so daß

$$f = a(t - \lambda_1) \cdot \dots \cdot (t - \lambda_n).$$

Da a der höchste Koeffizient ist, gilt $a \in \mathbb{R}$.

Seien $\lambda_1, \dots, \lambda_k \in \mathbb{R}$ und $\lambda_{k+1}, \dots, \lambda_n \in \mathbb{C}$. Nach dem Lemma ist $n - k$, d.h. die mit Vielfachheit gezählte Anzahl der nicht reellen Nullstellen, gerade, und durch Umnummerierung kann man erreichen, daß

$$\lambda_{k+1} = \bar{\lambda}_{k+2}, \dots, \lambda_{n-1} = \bar{\lambda}_n$$

gilt. Jedes Paar $\lambda, \bar{\lambda}$ konjugiert komplexer Nullstellen mit $\lambda = \alpha + i\beta$ kann man also (wie wir im Hilfssatz gesehen haben) zu einem normierten quadratischen Faktor

$$g = (t - \lambda)(t - \bar{\lambda}) = t^2 - 2\alpha t + (\alpha^2 + \beta^2) \in \mathbb{R}[t]$$

zusammenfassen. g hat keine reelle Nullstelle, was man auch an der *Diskriminante* ablesen kann:

$$4\alpha^2 - 4(\alpha^2 + \beta^2) = -4\beta^2 < 0.$$

Damit ist folgendes Ergebnis bewiesen:

Theorem. Jedes Polynom $f \in \mathbb{R}[t]$ mit $\deg f = n \geq 1$ gestattet eine Zerlegung

$$f = a(t - \lambda_1) \cdot \dots \cdot (t - \lambda_r) \cdot g_1 \cdot \dots \cdot g_m,$$

wobei $a, \lambda_1, \dots, \lambda_r$ reell sind, mit $a \neq 0$, und $g_1, \dots, g_m \in \mathbb{R}[t]$ normierte Polynome vom Grad 2 ohne reelle Nullstellen sind. Insbesondere ist $n = r + 2m$.

Korollar. Jedes Polynom $f \in \mathbb{R}[t]$ von ungeradem Grad hat mindestens eine reelle Nullstelle.

Dies folgt sofort aus der Gleichung $n = r + 2m$. Natürlich kann man die Aussage des Korollars mit Hilfe des Zwischenwertsatzes viel einfacher direkt beweisen.

Es sei erwähnt, daß man den Fundamentalsatz der Algebra ohne weitere Hilfsmittel der Analysis aus diesem Korollar ableiten kann (vgl. etwa [F-S]).

Der Fundamentalsatz der Algebra ist eine reine Existenzaussage, d.h. man erhält daraus kein Verfahren zur praktischen Bestimmung der Nullstellen. Für ein Polynom

$$at^2 + bt + c \in \mathbb{C}[t]$$

vom Grad 2 kann man die Nullstelle nach der Formel

$$\frac{-b \pm \sqrt{b^2 - 4ac}}{2a}$$

berechnen. Etwas kompliziertere Formeln dieser Art gibt es auch für die Nullstellen von Polynomen vom Grad 3 und 4.

Wie erstmals N.H. ABEL im Jahre 1826 zeigte, kann es solche allgemeine Formeln für Polynome vom Grad größer als 4 aus algebraischen Gründen nicht geben. Man ist daher weitgehend auf Näherungsverfahren zur Approximation der Nullstellen angewiesen.

1.3.11. Wie gerade erläutert, gibt es keine allgemein gültige Formel zur Berechnung der Nullstellen eines Polynoms aus den Koeffizienten. Die umgekehrte Aufgabe ist jedoch ganz einfach: Ist $f \in K[t]$ und

$$f(t) = t^n + \alpha_{n-1}t^{n-1} + \dots + \alpha_1t + \alpha_0 = (t - \lambda_1) \cdot \dots \cdot (t - \lambda_n),$$

d.h. zerfällt f in Linearfaktoren, so folgt durch ausmultiplizieren der rechten Seite

$$\begin{aligned}\alpha_0 &= (-1)^n \lambda_1 \cdot \dots \cdot \lambda_n, \\ \alpha_1 &= (-1)^{n-1} (\lambda_2 \cdot \dots \cdot \lambda_n + \lambda_1 \lambda_3 \cdot \dots \cdot \lambda_n + \dots + \lambda_1 \cdot \dots \cdot \lambda_{n-1}), \\ &\vdots \\ \alpha_{n-2} &= \lambda_1 \lambda_2 + \dots + \lambda_1 \lambda_n + \lambda_2 \lambda_3 + \dots + \lambda_2 \lambda_n + \dots + \lambda_{n-1} \lambda_n, \\ \alpha_{n-1} &= -(\lambda_1 + \dots + \lambda_n).\end{aligned}$$

Um das formal besser aufschreiben zu können, definieren wir für $k = 1, \dots, n$ die *elementarsymmetrischen Funktionen*

$$s_k(\lambda_1, \dots, \lambda_n) := \sum_{1 \leq i_1 < \dots < i_k \leq n} \lambda_{i_1} \cdot \dots \cdot \lambda_{i_k}.$$

Die Summe hat so viele Summanden, wie es Möglichkeiten gibt, k verschiedene Indizes $i_1, \dots, i_k \in \{1, \dots, n\}$ auszuwählen, also $\binom{n}{k}$ Summanden. Für die Koeffizienten von f gilt dann

$$\alpha_k = (-1)^{n-k} s_{n-k}(\lambda_1, \dots, \lambda_n).$$

Diese Aussage nennt man den *Wurzelsatz von VIETA*. Er gibt an, wie sich die Koeffizienten aus den Nullstellen berechnen.

Wie schon gesagt, ist umgekehrt die Bestimmung der Nullstellen aus den Koeffizienten im Allgemeinen ein sehr schwieriges Problem. Im Spezialfall eines reellen Polynoms gibt es immerhin einen Zusammenhang zwischen den Vorzeichen der Koeffizienten und den Nullstellen. Die Tendenz ist klar: Die Vorzeichen der Koeffizienten beeinflussen den Verlauf der Werte und damit die Lage der reellen Nullstellen eines Polynoms.

Vorweg behandeln wir den elementarsten Fall, der später (in 5.7.3) beim Test der Definitheit einer Matrix nützlich sein wird. Wir machen dabei die sehr einschränkende Voraussetzung, daß das Polynom $f \in \mathbb{R}[t]$ (mit Vielfachheit gezählt) so viele reelle Nullstellen hat, wie sein Grad angibt.

Vorzeichenregel. *Angenommen, das reelle Polynom*

$$f(t) = t^n + \alpha_{n-1}t^{n-1} + \dots + \alpha_1t + \alpha_0$$

hat reelle Nullstellen $\lambda_1, \dots, \lambda_n$. Dann gilt:

- a) *Genau dann sind alle Nullstellen λ_i negativ, wenn alle Koeffizienten α_j positiv sind.*

b) Genau dann sind alle Nullstellen λ_i positiv, wenn die Vorzeichen der Koeffizienten α_j alternierend sind, d.h. es ist

$$(-1)^{n-j} \alpha_j > 0 \text{ für } j = 0, \dots, n-1.$$

Beweis. Es genügt, den Fall a) zu beweisen; Fall b) folgt daraus sofort, indem man das Polynom f_- mit $f_-(t) := f(-t)$ betrachtet.

Sind alle $\alpha_j > 0$ und ist $\lambda \geq 0$, so ist

$$f(\lambda) = \lambda^n + \alpha_{n-1} \lambda^{n-1} + \dots + \alpha_0 \geq \alpha_0 > 0.$$

Also sind alle Nullstellen negativ.

Sind alle $\lambda_i < 0$, so ist s_k positiv für gerades k und negativ für ungerades k . Also ist

$$\alpha_j = (-1)^{n-j} \cdot s_{n-j} > 0$$

für alle j . □

Um eine allgemeinere Aussage zu erhalten, definiert man zunächst für ein beliebiges Polynom

$$g(t) := t^m + \beta_{m-1} t^{m-1} + \dots + \beta_0 \in \mathbb{R}[t]$$

die Zahl $Z(g)$ der *Zeichenwechsel* von g wie folgt: Man schreibt die Koeffizienten

$$1, \beta_{m-1}, \beta_{m-2}, \dots, \beta_1, \beta_0$$

auf und zählt von links beginnend ab, wie oft die Vorzeichen der β_i wechseln. Koeffizienten Null werden dabei einfach übergangen. Für den gleich folgenden Beweis schreiben wir das noch etwas präziser auf. Es gibt eindeutig bestimmte Indizes k_p mit

$$m > k_1 > k_2 > \dots > k_r \geq 0$$

derart, daß k_1 (von oben beginnend) der erste Index mit $\beta_{k_1} < 0$ ist, k_2 der erste darauf folgende mit $\beta_{k_2} > 0$ usw. Schließlich findet bei β_{k_r} der letzte Vorzeichenwechsel statt. Dann ist offensichtlich

$$\beta_{k_p} \begin{cases} > 0 \text{ für } p \text{ gerade,} \\ < 0 \text{ für } p \text{ ungerade,} \end{cases} \quad \text{und } r = Z(g) \leq m = \deg g.$$

Ohne Voraussetzung über die Gesamtzahl der reellen Nullstellen eines Polynoms $f \in \mathbb{R}[t]$ kann man nun mit Hilfe der Vorzeichen der Koeffizienten eine Abschätzung für die Zahlen

$N_+(f) := \text{Anzahl der positiven Nullstellen von } f$ und

$N_-(f) := \text{Anzahl der negativen Nullstellen von } f$

geben. Dabei sei wie oben $f_- \in \mathbb{R}[t]$ definiert durch $f_-(t) := f(-t)$.

Vorzeichenregel von DESCARTES. Für ein Polynom

$$f(t) = t^n + \alpha_{n-1} t^{n-1} + \dots + \alpha_0 \in \mathbb{R}[t]$$

mit $\alpha_0 \neq 0$ gilt

$$N_+(f) \leq Z(f) \text{ und } N_-(f) \leq Z(f_-).$$

Mit etwas zusätzlichem Aufwand kann man zeigen, daß die Differenz

$$Z(f) - N_+(f) \text{ immer gerade ist.}$$

Beweis. Da beim Übergang von f zu f_- die Vorzeichen der Nullstellen umgekehrt werden, genügt es, die erste Ungleichung zu zeigen. Dazu benutzen wir das

Lemma. Ist $g(t) = t^m + \beta_{m-1}t^{m-1} + \dots + \beta_0 \in \mathbb{R}[t]$ mit $\beta_0 \neq 0$ und $0 < \lambda \in \mathbb{R}$, so folgt

$$Z((t - \lambda) \cdot g) \geq Z(g) + 1.$$

Bei Multiplikation mit einem Linearfaktor mit positiver Nullstelle nehmen also die Zeichenwechsel zu.

Sind $\lambda_1, \dots, \lambda_k$ mit $k = N_+(f)$ die positiven reellen Nullstellen von f , so gibt es eine Zerlegung

$$f(t) = (t - \lambda_1) \cdot \dots \cdot (t - \lambda_k) \cdot g(t)$$

mit $g(t) \in \mathbb{R}[t]$. Wegen $Z(g) \geq 0$ folgt durch k -malige Anwendung des Lemmas $Z(f) \geq Z(g) + k \geq k = N_+(f)$. $\square$

Es bleibt der etwas knifflige *Beweis des Lemmas*. Für die Koeffizienten des Produktpolynoms

$$(t - \lambda)(t^m + \beta_{m-1}t^{m-1} + \dots + \beta_0) = t^{m+1} + \gamma_m t^m + \dots + \gamma_0$$

gilt:

$$\gamma_m = \beta_{m-1} - \lambda,$$

$$\gamma_j = \beta_{j-1} - \lambda \beta_j, \quad \text{für } j = 1, \dots, m-1,$$

$$\gamma_0 = -\lambda \beta_0.$$

Wir betrachten die oben angegebenen kritischen Indizes $k_1, \dots, k_r$. Ist ρ ungerade, so folgt

$$\beta_{k_\rho} < 0, \quad \beta_{k_\rho+1} \geq 0, \quad \text{also } \gamma_{k_\rho+1} = \beta_{k_\rho} - \lambda \beta_{k_\rho+1} < 0.$$

Für gerade ρ ist

$$\beta_{k_\rho} > 0, \quad \beta_{k_\rho+1} \leq 0, \quad \text{also } \gamma_{k_\rho+1} = \beta_{k_\rho} - \lambda \beta_{k_\rho+1} > 0.$$

Das bedeutet, daß β_{k_ρ} und $\gamma_{k_\rho+1}$ für $\rho = 1, \dots, r$ dasselbe Vorzeichen haben.

Wegen $\gamma_0 = -\lambda\beta_0$ haben γ_0 und β_0 entgegengesetzte Vorzeichen. Daher hat $(t - \lambda) \cdot g$ mindestens einen Vorzeichenwechsel mehr als g . $\square$

Noch allgemeiner kann man mit Hilfe so genannter STURMScher Ketten die Anzahl der Nullstellen abschätzen, die in einem vorgegebenen Intervall liegen (vgl. dazu etwa [D], [St], [Wi]).

Aufgaben zu 1.3

1. Bestimmen Sie (bis auf Isomorphie) alle Körper mit 3 bzw. 4 Elementen.

2. K und K' seien zwei Körper und $\varphi: K \rightarrow K'$ ein Ringhomomorphismus. Zeigen Sie, dass φ entweder injektiv oder der Nullhomomorphismus ist.

3. Ist R ein Ring, M eine beliebige nichtleere Menge und $S = \text{Abb}(M; R)$ die Menge aller Abbildungen von M nach R , so ist auf S durch

$$(f + g)(m) := f(m) + g(m), (f \cdot g)(m) := f(m) \cdot g(m),$$

eine Addition und eine Multiplikation erklärt.

a) Zeigen Sie, dass S auf diese Weise zu einem Ring wird.

b) Ist S ein Körper, falls R ein Körper ist?

4.* Sei $p \in \mathbb{N}$ eine Primzahl und $n \in \mathbb{N} \setminus \{0\}$. Zeigen Sie, dass es einen Körper mit p^n Elementen gibt.

5. Sei K' ein Körper und K ein Unterkörper von K' .

Zeigen Sie: Sind $f, g \in K[t]$, $q \in K'[t]$ mit $f = qg$, so folgt bereits $q \in K[t]$.

6. Sei K ein Körper und $x_0, \dots, x_n, y_0, \dots, y_n \in K$ mit $x_i \neq x_j$ für alle $i \neq j$. Zeigen Sie, dass es genau ein Polynom $f \in K[t]$ vom Grad $\leq n$ gibt, so dass $f(x_i) = y_i$ für $i = 0, \dots, n$. Hinweis: Konstruieren Sie zuerst Polynome $g_k \in K[t]$ vom Grad $\leq n$ mit

$$g_k(x_i) = \begin{cases} 1 & \text{für } i = k, \\ 0 & \text{für } i \neq k. \end{cases}$$

7. Seien $f, g \in \mathbb{C}[t]$ Polynome mit $\mu(f, \lambda) \leq \mu(g, \lambda)$ für alle $\lambda \in \mathbb{C}$. Zeigen Sie, dass dann f ein Teiler von g ist. Gilt diese Aussage auch in $\mathbb{R}[t]$?

8. Sei K ein Körper und $\sim: K[t] \rightarrow \text{Abb}(K, K)$, $f \mapsto \tilde{f}$, die Abbildung aus 1.3.5, die jedem Polynom f die zugehörige Abbildung $\tilde{f}$ zuordnet. Zeigen Sie, dass $\sim$ surjektiv, aber nicht injektiv ist, falls der Körper K endlich ist.

9. Analog zu 1.3.5 definiert man ein Polynom mit Koeffizienten über einem Körper K in n Unbestimmten $t_1, \dots, t_n$ als einen formalen Ausdruck der Gestalt

$$f(t_1, \dots, t_n) = \sum_{0 \leq i_1, \dots, i_n \leq k} a_{i_1 \dots i_n} \cdot t_1^{i_1} \cdot \dots \cdot t_n^{i_n},$$

wobei $k \in \mathbb{N}$ und $a_{i_1 \dots i_n} \in K$. $K[t_1, \dots, t_n]$ bezeichne die Menge all solcher Polynome. Wie für Polynome in einer Unbestimmten kann auch in $K[t_1, \dots, t_n]$ eine Addition und eine Multiplikation erklärt werden. Sind $f, g \in K[t_1, \dots, t_n]$, so erfolgt die Addition von f und g koeffizientenweise und die Multiplikation wieder durch formales Ausmultiplizieren.

- a) Finden Sie Formeln für die Addition und Multiplikation von Polynomen in $K[t_1, \dots, t_n]$, und zeigen Sie, dass $K[t_1, \dots, t_n]$ auf diese Weise zu einem nullteilerfreien, kommutativen Ring wird.

Ein Polynom $h \in K[t_1, \dots, t_n] \setminus \{0\}$ heißt *homogen* (vom Grad d), falls

$$h = \sum_{i_1 + \dots + i_n = d} a_{i_1 \dots i_n} \cdot t_1^{i_1} \cdot \dots \cdot t_n^{i_n}.$$

- b) Für ein homogenes Polynom $h \in K[t_1, \dots, t_n]$ vom Grad d gilt:

$$h(\lambda t_1, \dots, \lambda t_n) = \lambda^d \cdot h(t_1, \dots, t_n) \quad \text{für alle } \lambda \in K.$$

- c) Ist K unendlich und $f \in K(t_1, \dots, t_n) \setminus \{0\}$, so folgt aus

$$f(\lambda t_1, \dots, \lambda t_n) = \lambda^d \cdot f(t_1, \dots, t_n) \quad \text{für alle } \lambda \in K,$$

dass f homogen vom Grad d ist.

- d) Ist h_1 homogen vom Grad d_1 und h_2 homogen vom Grad d_2 , so ist $h_1 \cdot h_2$ homogen vom Grad $d_1 + d_2$.

10. Sei K ein Körper und $K[t]$ der Polynomring in einer Unbestimmten.

- a) Zeigen Sie, dass in der Menge $K[t] \times (K[t] \setminus \{0\})$ durch

$$(g, h) \sim (g', h') \Leftrightarrow gh' = g'h$$

eine Äquivalenzrelation gegeben ist.

$K(t)$ sei die Menge der Äquivalenzklassen. Die zu (g, h) gehörige Äquivalenzklasse sei mit $\frac{g}{h}$ bezeichnet. Somit ist $\frac{g}{h} = \frac{g'}{h'} \Leftrightarrow gh' = g'h$.

- b) Zeigen Sie, dass in $K(t)$ die Verknüpfungen

$$\frac{g}{h} + \frac{g'}{h'} := \frac{gh' + hg'}{hh'}, \quad \frac{g}{h} \cdot \frac{g'}{h'} := \frac{gg'}{hg'},$$

wohl definiert sind (vgl. 1.2.7).

- c) Zeigen Sie schließlich, dass $K(t)$ mit diesen Verknüpfungen zu einem Körper wird. Man nennt $K(t)$ den Körper der rationalen Funktionen.

11. Was folgt aus der Vorzeichenregel von DESCARTES für das Polynom $t^n + 1$?

12. Man folgere die spezielle Vorzeichenregel aus der Vorzeichenregel von DESCARTES.

1.4 Vektorräume

In diesem ganzen Abschnitt bezeichnet K einen Körper. Wer die vorhergehenden Definitionen übersprungen hat, kann sich zunächst mit dem äußerst wichtigen Spezialfall $K = \mathbb{R}$ begnügen.

1.4.1. Bevor wir den allgemeinen Begriff des Vektorraums einführen, einige

Beispiele. a) Das Standardbeispiel ist der *Standardraum*

$$K^n = \{x = (x_1, \dots, x_n) : x_i \in K\}.$$

Mit Hilfe der Addition und Multiplikation in K erhält man zwei neue Verknüpfungen

$$\begin{aligned} \dot{+} : K^n \times K^n &\rightarrow K^n, & (x, y) &\mapsto x \dot{+} y, & \text{ und} \\ \cdot : K \times K^n &\rightarrow K^n, & (\lambda, x) &\mapsto \lambda \cdot x, \end{aligned}$$

durch

$$\begin{aligned} (x_1, \dots, x_n) \dot{+} (y_1, \dots, y_n) &:= (x_1 + y_1, \dots, x_n + y_n) \quad \text{ und} \\ \lambda \cdot (x_1, \dots, x_n) &:= (\lambda x_1, \dots, \lambda x_n). \end{aligned}$$

Zur vorübergehenden Unterscheidung sind die neuen Verknüpfungen mit $\dot{+}$ und $\cdot$ bezeichnet. In K wird die Addition durch $+$ und die Multiplikation ohne Symbol ausgedrückt. Man beachte, daß nur $\dot{+}$ eine Verknüpfung im Sinn von 1.2.1 ist (solche Verknüpfungen nennt man manchmal auch *innere* im Gegensatz zur *äußeren* $\cdot$).

b) In der Menge $M(m \times n; K)$ der Matrizen mit m Zeilen, n Spalten und Einträgen aus K kann man addieren und mit Skalaren multiplizieren:

Ist $A = (a_{ij})$, $B = (b_{ij}) \in M(m \times n; K)$ und $\lambda \in K$, so sind

$$A \dot{+} B := (a_{ij} + b_{ij}) \quad \text{ und } \quad \lambda \cdot A := (\lambda a_{ij}) \in M(m \times n; K).$$

Bei der Addition werden also die Einträge an den entsprechenden Stellen addiert, bei Multiplikation mit λ alle Einträge gleich multipliziert. Bis auf die andere Art, die Einträge aufzuschreiben, ist dieses Beispiel gleich $K^{m \cdot n}$ aus a).

c) Im Körper $\mathbb{C}$ der komplexen Zahlen kann man mit reellen Zahlen multiplizieren, das ergibt eine Abbildung

$$\mathbb{R} \times \mathbb{C} \rightarrow \mathbb{C}, \quad (\lambda, a + ib) \mapsto \lambda a + i\lambda b.$$

d) Im Polynomring $K[t]$ kann man neben Addition und Multiplikation von Polynomen eine weitere Multiplikation

$$\cdot : K \times K[t] \rightarrow K[t], \quad (\lambda, f) \mapsto \lambda \cdot f,$$

mit Elementen aus K erklären durch

$$\lambda \cdot (a_0 + a_1 t + \dots + a_n t^n) := \lambda a_0 + (\lambda a_1) t + \dots + (\lambda a_n) t^n.$$

e) Ist M eine beliebige Menge und

$$V = \{f: M \rightarrow K\} = \text{Abb}(M, K)$$

die Menge aller Abbildungen, so sind für $f, g \in V$ und $\lambda \in K$

$$f \dot{+} g \in V \quad \text{und} \quad \lambda \cdot f \in V$$

erklärt durch

$$(f \dot{+} g)(x) := f(x) + g(x) \quad \text{und} \quad (\lambda \cdot f)(x) := \lambda f(x).$$

Offensichtlich erhält man im Spezialfall $M = \{1, \dots, n\}$ wieder das Standardbeispiel a).

Als Extrakt aus diesen Beispielen führen wir nun den wichtigsten Begriff der linearen Algebra ein (zur Entstehung vgl. [Koe], Kap. 1, §2):

Definition. Sei K ein Körper. Eine Menge V zusammen mit einer inneren Verknüpfung

$$\dot{+}: V \times V \rightarrow V, \quad (v, w) \mapsto v \dot{+} w,$$

(Addition genannt) und einer äußeren Verknüpfung

$$\cdot: K \times V \rightarrow V, \quad (\lambda, v) \mapsto \lambda \cdot v,$$

(Multiplikation mit Skalaren oder skalare Multiplikation genannt) heißt K -Vektorraum (oder Vektorraum über K), wenn folgendes gilt:

V1 V zusammen mit der Addition ist eine abelsche Gruppe (das neutrale Element heißt Nullvektor, es wird mit 0 , und das Negative wird mit $-v$ bezeichnet).

V2 Die Multiplikation mit Skalaren muß in folgender Weise mit den anderen Verknüpfungen verträglich sein:

$$(\lambda + \mu) \cdot v = \lambda \cdot v \dot{+} \mu \cdot v, \quad \lambda \cdot (v \dot{+} w) = \lambda \cdot v \dot{+} \lambda \cdot w,$$

$$\lambda \cdot (\mu \cdot v) = (\lambda\mu) \cdot v, \quad 1 \cdot v = v,$$

für alle $\lambda, \mu \in K$ und $v, w \in V$.

Man beachte, daß, wie in Beispiel a) erläutert, die Verknüpfungen in K und V vorübergehend verschieden bezeichnet werden.

Durch Einsetzen der Definitionen und elementarste Rechnungen sieht man, daß in den obigen Beispielen a) bis e) die Vektorraumaxiome erfüllt sind. Dabei ist in Beispiel a) der Nullvektor gegeben durch $0 = (0, \dots, 0)$ und das Negative durch

$$-(x_1, \dots, x_n) = (-x_1, \dots, -x_n).$$

In Beispiel c) wird $\mathbb{C}$ zu einem $\mathbb{R}$ -Vektorraum.

Bemerkung. In einem K -Vektorraum V hat man die folgenden weiteren Rechenregeln:

a) $0 \cdot v = 0$.

b) $\lambda \cdot 0 = 0$.

c) $\lambda \cdot v = 0 \Rightarrow \lambda = 0 \text{ oder } v = 0$.

d) $(-1) \cdot v = -v$.

Beweis. a) $0 \cdot v = (0 + 0) \cdot v = 0 \cdot v + 0 \cdot v$.

b) $\lambda \cdot 0 = \lambda \cdot (0 + 0) = \lambda \cdot 0 + \lambda \cdot 0$.

c) Ist $\lambda \cdot v = 0$, aber $\lambda \neq 0$, so folgt

$$v = 1 \cdot v = (\lambda^{-1} \lambda) \cdot v = \lambda^{-1} \cdot (\lambda \cdot v) = \lambda^{-1} \cdot 0 = 0.$$

d) $v + (-1) \cdot v = 1 \cdot v + (-1) \cdot v = (1 - 1) \cdot v = 0 \cdot v = 0$. □

Die Axiome und die daraus abgeleiteten Regeln zeigen insbesondere, daß es völlig ungefährlich ist, wenn man die Verknüpfungen in K und V gleich bezeichnet und auch den Nullvektor 0 abmagert zu 0 . Das wollen wir ab sofort tun. Was gemeint ist, wird jeweils aus dem Zusammenhang klar werden.

1.4.2. In Kapitel 0 hatten wir homogene lineare Gleichungssysteme der Form $A \cdot x = 0$ mit reellen Koeffizienten betrachtet. Die Lösungen sind Teilmengen $W \subset \mathbb{R}^n$, ihre präzise Beschreibung ist unser Ziel. Schlüssel dafür ist die Beobachtung, daß W von $\mathbb{R}^n$ eine Vektorraumstruktur erbt. Allgemeiner ist die Frage, wann das für eine Teilmenge $W \subset V$ eines Vektorraumes der Fall ist.

Definition. Sei V ein K -Vektorraum und $W \subset V$ eine Teilmenge. W heißt *Untervektorraum von V* , falls folgendes gilt:

UV1 $W \neq \emptyset$.

UV2 $v, w \in W \Rightarrow v + w \in W$
(d.h. W ist abgeschlossen gegenüber der Addition).

UV3 $v \in W, \lambda \in K \Rightarrow \lambda v \in W$
(d.h. W ist abgeschlossen gegenüber der Multiplikation mit Skalaren).

Beispiele. a) In $V = \mathbb{R}^2$ betrachten wir die Teilmengen

$$\begin{aligned} W_1 &= \{0\}, \\ W_2 &= \{(x_1, x_2) \in \mathbb{R}^2 : a_1 x_1 + a_2 x_2 = b\}, \\ W_3 &= \{(x_1, x_2) \in \mathbb{R}^2 : x_1^2 + x_2^2 \leq 1\}, \\ W_4 &= \{(x_1, x_2) \in \mathbb{R}^2 : x_1 \geq 0, x_2 \geq 0\}, \\ W_5 &= \{(x_1, x_2) \in \mathbb{R}^2 : x_1 \cdot x_2 \geq 0\}. \end{aligned}$$

Der Leser prüfe die Bedingungen UV2 und UV3 nach. Nur für W_1 und für W_2 mit $b = 0$ sind beide erfüllt.

b) Ist A eine reelle $m \times n$ -Matrix, so ist die Lösungsmenge

$$W := \{x \in \mathbb{R}^n : Ax = 0\}$$

des zugehörigen homogenen linearen Gleichungssystems (vgl. 0.4.2) ein Untervektorraum des $\mathbb{R}^n$. Das rechnet man ganz leicht nach.

c) Im Vektorraum $V = \text{Abb}(\mathbb{R}, \mathbb{R})$ (vgl. 1.4.1, Beispiel e)) hat man die Untervektorräume

$$\mathbb{R}[t]_d \subset \mathbb{R}[t] \subset \mathcal{D}(\mathbb{R}, \mathbb{R}) \subset \mathcal{C}(\mathbb{R}, \mathbb{R}) \subset \text{Abb}(\mathbb{R}, \mathbb{R})$$

der Polynome vom Grad $\leq d$, aller Polynome, der differenzierbaren Funktionen und der stetigen Funktionen.

1.4.3. Aus den Eigenschaften UV2 und UV3 folgt, daß Addition und Multiplikation mit Skalaren von V auf W induziert werden.

Satz. Ein Untervektorraum $W \subset V$ ist zusammen mit der induzierten Addition und Multiplikation mit Skalaren wieder ein Vektorraum.

Mit Hilfe dieses Satzes kann man sich in vielen Fällen (etwa Beispiel b) und c) aus 1.4.2) den langweiligen Nachweis aller Vektorraumaxiome für W sparen, wenn man das ein für alle mal schon in einem größeren V getan hatte.

Beweis. Die Eigenschaften V2 sowie Kommutativ- und Assoziativgesetz der Addition gelten in W , da sie in V gelten. Der Nullvektor 0 liegt in W , da wegen UV1 ein $v \in W$ existiert, woraus $0 = 0 \cdot v \in W$ mit UV3 folgt. Zu jedem $v \in W$ ist wegen UV3 auch $-v = (-1) \cdot v \in W$. $\square$

Noch eine kleine Pflichtübung zur Erzeugung neuer Untervektorräume aus alten:

Lemma. Sei V ein Vektorraum, I eine beliebige Indexmenge, und für jedes $i \in I$ sei ein Untervektorraum W_i gegeben. Dann ist der Durchschnitt

$$W := \bigcap_{i \in I} W_i \subset V$$

wieder ein Untervektorraum.

Beweis. Da 0 in allen W_i enthalten ist, ist auch $0 \in W$, also $W \neq \emptyset$. Sind $v, w \in W$, so sind v, w in allen W_i enthalten. Da dann auch $v + w$ in allen W_i enthalten ist, ist $v + w$ in W enthalten. Ganz analog beweist man, daß mit $\lambda \in K$ und $v \in W$ auch $\lambda v \in W$ gilt. $\square$

Beispiel. Ist $V = \mathbb{R}[t]$ der Vektorraum aller reellen Polynome, $I = \mathbb{N}$ und $W_d := \mathbb{R}[t]_d$ der Untervektorraum der Polynome vom Grad $\leq d$, so ist

$$\bigcap_{d \in \mathbb{N}} W_d = W_0 = \mathbb{R}.$$

Vorsicht! Die Vereinigung von Untervektorräumen ist im allgemeinen kein Untervektorraum. Man mache sich das an Beispielen klar (etwa zwei Geraden im $\mathbb{R}^2$). Es gilt sogar die

Bemerkung. Sind $W, W' \subset V$ Untervektorräume derart, daß $W \cup W'$ Untervektorraum ist, so ist $W \subset W'$ oder $W' \subset W$.

Beweis. Angenommen, es ist $W \not\subset W'$. Dann ist $W' \subset W$ zu zeigen. Ist $w' \in W'$ und $w \in W \setminus W'$, so sind $w, w' \in W \cup W'$, also auch $w + w' \in W \cup W'$. $w + w'$ kann nicht Element von W' sein, denn sonst wäre $w = w + w' - w' \in W'$. Also ist $w + w' \in W$ und auch $w' = w + w' - w \in W$. $\square$

1.4.4. Eine Teilmenge eines Vektorraumes, die kein Untervektorraum ist, kann man zu einem solchen *abschließen*. Wie das geht, wird nun beschrieben.

Wir betrachten eine noch etwas allgemeinere Situation, nämlich einen K -Vektorraum V und eine Familie $(v_i)_{i \in I}$ von Vektoren $v_i \in V$ (vgl. 1.1.6). Ist $I = \{1, \dots, r\}$, so hat man Vektoren $v_1, \dots, v_r$. Ein $v \in V$ heißt *Linearkombination* von $v_1, \dots, v_r$, wenn es $\lambda_1, \dots, \lambda_r \in K$ gibt, so daß

$$v = \lambda_1 v_1 + \dots + \lambda_r v_r.$$

Für allgemeines I definiert man

$$\text{span}_K(v_i)_{i \in I}$$

als die Menge all der $v \in V$, die sich aus einer (von v abhängigen) endlichen Teilfamilie von $(v_i)_{i \in I}$ linear kombinieren lassen. Um das präzise aufschreiben

zu können, benötigt man Doppelindizes: Zu $v \in V$ muß es eine Zahl $r \in \mathbb{N}$ sowie Indizes $i_1, \dots, i_r \in I$ und Skalare $\lambda_1, \dots, \lambda_r$ geben, so daß

$$v = \lambda_1 v_{i_1} + \dots + \lambda_r v_{i_r}.$$

Man nennt $\text{span}_K(v_i)_{i \in I}$ den von der Familie *aufgespannten* (oder *erzeugten*) Raum. Ist $I = \emptyset$, so setzt man

$$\text{span}_K(v_i)_{i \in \emptyset} := \{0\}.$$

Für eine endliche Familie $(v_1, \dots, v_r)$ verwendet man oft die suggestivere Notation

$$\begin{aligned} K v_1 + \dots + K v_r &:= \text{span}_K(v_1, \dots, v_r) \\ &= \{v \in V : \text{es gibt } \lambda_1, \dots, \lambda_r \in K \text{ mit } v = \lambda_1 v_1 + \dots + \lambda_r v_r\}. \end{aligned}$$

Falls klar ist, welcher Körper gemeint ist, schreibt man nur span statt span_K .

Bemerkung. Sei V ein K -Vektorraum und $(v_i)_{i \in I}$ eine Familie von Elementen aus V . Dann gilt:

- a) $\text{span}(v_i) \subset V$ ist Untervektorraum.
- b) Ist $W \subset V$ Untervektorraum, und gilt $v_i \in W$ für alle $i \in I$, so ist $\text{span}(v_i) \subset W$.

Kurz ausgedrückt: $\text{span}(v_i)$ ist der *kleinste* Untervektorraum von V , der alle v_i enthält.

Beweis. a) ist ganz klar. Sind alle v_i in W enthalten, so sind auch alle endlichen Linearkombinationen aus den v_i in W enthalten, denn W ist Untervektorraum. Daraus folgt b). $\square$

Ist $M \subset V$ eine Teilmenge, so ist entsprechend $\text{span}(M)$ erklärt als die Menge aller endlichen Linearkombinationen von Vektoren aus M , und das ist der kleinste Untervektorraum mit

$$M \subset \text{span}(M) \subset V.$$

Es mag auf den ersten Blick nicht recht einleuchten, warum man bei der Erzeugung eines Untervektorraumes allgemeiner von einer Familie von Vektoren ausgeht. Das hat den Vorteil, daß es bei einer Familie (im Gegensatz zu einer Menge) sinnvoll ist, wenn man sagt „ein Vektor kommt mehrfach vor“. Ist $I = \{1, \dots, n\}$, so haben die Vektoren der Familie außerdem eine natürliche Reihenfolge.

Beispiele. a) Sei $V = \mathbb{R}^3$. Sind $v_1, v_2 \in \mathbb{R}^3$, so ist $\text{span}(v_1)$ die Gerade durch 0 und v_1 , wenn $v_1 \neq 0$. $\text{span}(v_1, v_2)$ ist die Ebene durch 0, v_1 und v_2 , falls $v_2 \notin \text{span}(v_1)$.

b) Im K^n mit $I = \{1, \dots, n\}$ erklären wir für $i \in I$

$$e_i := (0, \dots, 0, 1, 0, \dots, 0),$$

wobei die 1 an der i -ten Stelle steht. Dann ist $\text{span}(e_i)_{i \in I} = K^n$.

c) Ist $V = K[t]$ der Polynomring, $I = \mathbb{N}$ und $v_n = t^n$, so ist

$$\text{span}(v_n)_{n \in \mathbb{N}} = K[t].$$

d) Betrachten wir in Beispiel a) aus 1.4.2 die W_j mit $j = 3, 4, 5$, so ist $\text{span}(W_j) = \mathbb{R}^2$.

1.4.5. Ein Untervektorraum kann von sehr vielen verschiedenen Familien erzeugt werden, und das kann mit unterschiedlicher Effizienz geschehen. Dazu betrachten wir die Beispiele aus 1.4.4. Bei b) ist die Situation optimal, denn es folgt für $x = (x_1, \dots, x_n) \in K^n$ und

$$x = \lambda_1 e_1 + \dots + \lambda_n e_n, \quad \text{daß} \quad \lambda_1 = x_1, \dots, \lambda_n = x_n.$$

Die Linearkombination ist also für jedes x eindeutig bestimmt, entsprechend in c). In den Beispielen aus d) hat man für jedes $x \in \mathbb{R}^2$ jeweils unendlich viele Möglichkeiten, es linear aus Elementen von W_j zu kombinieren.

Die Eindeutigkeit ist besonders einfach beim Nullvektor zu überprüfen. Bei beliebigen $v_1, \dots, v_r$ hat man die *triviale Linearkombination*

$$0 = 0v_1 + \dots + 0v_r.$$

Gibt es eine andere Linearkombination des Nullvektors, so ist die Eindeutigkeit der Darstellung verletzt. Das motiviert die

Definition. Sei V ein K -Vektorraum. Eine endliche Familie $(v_1, \dots, v_r)$ von Vektoren aus V heißt *linear unabhängig*, falls gilt: Sind $\lambda_1, \dots, \lambda_r \in K$ und ist

$$\lambda_1 v_1 + \dots + \lambda_r v_r = 0,$$

so folgt

$$\lambda_1 = \dots = \lambda_r = 0.$$

Anders ausgedrückt bedeutet das, daß sich der Nullvektor nur trivial aus den $v_1, \dots, v_r$ linear kombinieren läßt.

Eine beliebige Familie $(v_i)_{i \in I}$ von Vektoren aus V heißt *linear unabhängig*, falls jede endliche Teilfamilie linear unabhängig ist.

Die Familie $(v_i)_{i \in I}$ heißt *linear abhängig*, falls sie nicht linear unabhängig ist, d.h. falls es eine endliche Teilfamilie $(v_{i_1}, \dots, v_{i_r})$ und $\lambda_1, \dots, \lambda_r \in K$ gibt, die nicht alle gleich Null sind, so daß

$$\lambda_1 v_{i_1} + \dots + \lambda_r v_{i_r} = 0.$$

Zur Bequemlichkeit sagt man meist anstatt „die Familie $(v_1, \dots, v_r)$ von Vektoren aus V ist linear (un-)abhängig“ einfacher „die Vektoren $v_1, \dots, v_r \in V$ sind linear (un-)abhängig“.

Es ist vorteilhaft, auch die leere Familie, die den Nullvektorraum aufspannt, linear unabhängig zu nennen.

Die Definition der linearen Unabhängigkeit ist grundlegend für die ganze lineare Algebra, aber man muß sich etwas daran gewöhnen. Was sie geometrisch bedeutet, sieht man sehr gut an der Bedingung aus Aufgabe 1 zu 0.3 für zwei Vektoren im $\mathbb{R}^n$.

In der Definition der linearen Unabhängigkeit spielt der Nullvektor scheinbar eine besondere Rolle. Daß dem nicht so ist, zeigt das

Lemma. Für eine Familie $(v_i)_{i \in I}$ von Vektoren eines K -Vektorraumes sind folgende Bedingungen äquivalent:

- i) (v_i) ist linear unabhängig.
- ii) Jeder Vektor $v \in \text{span}(v_i)$ läßt sich in eindeutiger Weise aus Vektoren der Familie (v_i) linear kombinieren.

Beweis. ii) $\Rightarrow$ i) ist klar, denn bei einer linear abhängigen Familie hat der Nullvektor verschiedene Darstellungen.

i) $\Rightarrow$ ii): Sei ein $v \in \text{span}(v_i)$ auf zwei Arten linear kombiniert, also

$$v = \sum_{i \in I} \lambda_i v_i = \sum_{i \in I} \mu_i v_i, \quad (*)$$

wobei in beiden Summen jeweils nur endlich viele der Skalare λ_i und μ_i von Null verschieden sind. Es gibt also eine endliche Teilmenge $J \subset I$, so daß für jedes $\lambda_i \neq 0$ oder $\mu_i \neq 0$ der Index in J enthalten ist. Aus (*) folgt dann

$$\sum_{i \in I} (\lambda_i - \mu_i) v_i = 0,$$

und wegen der vorausgesetzten linearen Unabhängigkeit folgt $\lambda_i = \mu_i$ für alle $i \in J$ und somit auch für alle $i \in I$, da ja die restlichen λ_i und μ_i ohnehin Null waren. Damit ist die Eindeutigkeit der Linearkombination bewiesen. $\square$

Beispiele. a) Im K^n sind die Vektoren $e_1, \dots, e_n$ linear unabhängig.

b) Ist $A = (a_{ij}) \in M(m \times n; K)$ eine Matrix in Zeilenstufenform (vgl. 0.4.3), so sind die ersten r Zeilen $v_1, \dots, v_r$ von A linear unabhängig. Ist nämlich

$$A = \begin{pmatrix} \boxed{a_{1j_1}} & & & & \\ & \boxed{a_{2j_2}} & & & \\ & & 0 & & \\ & & & \ddots & \\ & & & & \boxed{a_{rj_r}} \end{pmatrix},$$

so folgt aus $\lambda_1 v_1 + \dots + \lambda_r v_r = 0$ zunächst $\lambda_1 a_{1j_1} = 0$, also $\lambda_1 = 0$ wegen $a_{1j_1} \neq 0$. Im zweiten Schritt folgt daraus analog $\lambda_2 = 0$ und weiter so $\lambda_3 = \dots = \lambda_r = 0$. Analog zeigt man, daß die Spalten von A mit den Indizes $j_1, j_2, \dots, j_r$ linear unabhängig sind.

c) Im Polynomring $K[t]$ ist die Familie $(t^n)_{n \in \mathbb{N}}$ linear unabhängig.

Weitere Beispiele finden sich in den Aufgaben 8 und 9. Noch ein paar weitere Kleinigkeiten:

Bemerkung. In jedem K -Vektorraum V gilt:

- a) Ein einziger Vektor $v \in V$ ist genau dann linear unabhängig, wenn $v \neq 0$.
- b) Gehört der Nullvektor zu einer Familie, so ist sie linear abhängig.
- c) Kommt der gleiche Vektor in einer Familie mehrmals vor, so ist sie linear abhängig.
- d) Ist $r \geq 2$, so sind die Vektoren $v_1, \dots, v_r$ genau dann linear abhängig, wenn einer davon Linearkombination der anderen ist.

Diese letzte Charakterisierung ist plausibler als die Definition, aber formal nicht so bequem zu handhaben.

Beweis. a) Ist v linear abhängig, so gibt es ein $\lambda \in K^*$ mit $\lambda v = 0$, also ist $v = 0$ nach Bemerkung c) in 1.4.1. Umgekehrt ist 0 linear abhängig, da $1 \cdot 0 = 0$.

b) $1 \cdot 0 = 0$.

c) Gibt es $i_1, i_2 \in I$ mit $i_1 \neq i_2$ aber $v_{i_1} = v_{i_2}$, so ist $1 \cdot v_{i_1} + (-1)v_{i_2} = 0$.

d) Sind die Vektoren $v_1, \dots, v_r$ linear abhängig, so gibt es $\lambda_1, \dots, \lambda_r \in K$ mit $\lambda_1 v_1 + \dots + \lambda_r v_r = 0$ und ein $k \in \{1, \dots, r\}$ mit $\lambda_k \neq 0$. Dann ist

$$v_k = -\frac{\lambda_1}{\lambda_k} v_1 - \dots - \frac{\lambda_{k-1}}{\lambda_k} v_{k-1} - \frac{\lambda_{k+1}}{\lambda_k} v_{k+1} - \dots - \frac{\lambda_r}{\lambda_k} v_r.$$

Ist umgekehrt $v_k = \mu_1 v_1 + \dots + \mu_{k-1} v_{k-1} + \mu_{k+1} v_{k+1} + \dots + \mu_r v_r$, so ist

$$\mu_1 v_1 + \dots + \mu_{k-1} v_{k-1} + (-1) v_k + \mu_{k+1} v_{k+1} + \dots + \mu_r v_r = 0. \quad \square$$

Aufgaben zu 1.4

1. Welche der folgenden Mengen sind Untervektorräume der angegebenen Vektorräume?

- a) $\{(x_1, x_2, x_3) \in \mathbb{R}^3 : x_1 = x_2 = 2x_3\} \subset \mathbb{R}^3$.
- b) $\{(x_1, x_2) \in \mathbb{R}^2 : x_1^2 + x_2^4 = 0\} \subset \mathbb{R}^2$.
- c) $\{(\mu + \lambda, \lambda^2) \in \mathbb{R}^2 : \mu, \lambda \in \mathbb{R}\} \subset \mathbb{R}^2$.
- d) $\{f \in \text{Abb}(\mathbb{R}, \mathbb{R}) : f(x) = f(-x) \text{ für alle } x \in \mathbb{R}\} \subset \text{Abb}(\mathbb{R}, \mathbb{R})$.
- e) $\{(x_1, x_2, x_3) \in \mathbb{R}^3 : x_1 \geq x_2\} \subset \mathbb{R}^3$.
- f) $\{A \in M(m \times n; \mathbb{R}) : A \text{ ist in Zeilenstufenform}\} \subset M(m \times n; \mathbb{R})$.

2. Seien V und W zwei K -Vektorräume. Zeigen Sie, dass das direkte Produkt $V \times W$ durch die Verknüpfungen

$$(v, w) + (v', w') := (v + v', w + w'), \quad \lambda \cdot (v, w) := (\lambda v, \lambda w),$$

ebenfalls zu einem K -Vektorraum wird.

3. Ist X eine nichtleere Menge, V ein K -Vektorraum und $\text{Abb}(X, V)$ die Menge aller Abbildungen von X nach V , so ist auf $\text{Abb}(X, V)$ durch

$$(f + g)(x) := f(x) + g(x), \quad (\lambda \cdot f)(x) := \lambda f(x),$$

eine Addition und eine skalare Multiplikation erklärt.

Zeigen Sie, dass $\text{Abb}(X, V)$ mit diesen Verknüpfungen zu einem K -Vektorraum wird.

4. Eine Abbildung $f: \mathbb{R} \rightarrow \mathbb{R}$ heißt 2π -periodisch, falls $f(x) = f(x + 2\pi)$ für alle $x \in \mathbb{R}$.

- a) Zeigen Sie, dass $V = \{f \in \text{Abb}(\mathbb{R}, \mathbb{R}) : f \text{ ist } 2\pi\text{-periodisch}\} \subset \text{Abb}(\mathbb{R}, \mathbb{R})$ ein Untervektorraum ist.
- b) Zeigen Sie, dass $W = \text{span}(\cos nx, \sin mx)_{n,m \in \mathbb{N}}$ ein Untervektorraum von V ist. (Man nennt W den Vektorraum der *trigonometrischen Polynome*.)

5. Seien

$$\ell^1 := \left\{ (x_i)_{i \in \mathbb{N}} : \sum_{i=0}^{\infty} |x_i| < \infty \right\} \subset \text{Abb}(\mathbb{N}, \mathbb{R}),$$

$$\ell^2 := \left\{ (x_i)_{i \in \mathbb{N}} : \sum_{i=0}^{\infty} |x_i|^2 < \infty \right\} \subset \text{Abb}(\mathbb{N}, \mathbb{R}),$$

$$\ell := \{ (x_i)_{i \in \mathbb{N}} : (x_i)_{i \in \mathbb{N}} \text{ konvergiert} \} \subset \text{Abb}(\mathbb{N}, \mathbb{R}),$$

$$\ell_{\infty} := \{ (x_i)_{i \in \mathbb{N}} : (x_i)_{i \in \mathbb{N}} \text{ beschränkt} \} \subset \text{Abb}(\mathbb{N}, \mathbb{R}).$$

Zeigen Sie, dass $\ell^1 \subset \ell^2 \subset \ell \subset \ell_{\infty} \subset \text{Abb}(\mathbb{N}, \mathbb{R})$ eine aufsteigende Kette von Untervektorräumen ist.

6. Kann eine abzählbar unendliche Menge M eine $\mathbb{R}$ -Vektorraumstruktur besitzen?

7. Gibt es eine $\mathbb{C}$ -Vektorraumstruktur auf $\mathbb{R}$, so dass die skalare Multiplikation $\mathbb{C} \times \mathbb{R} \rightarrow \mathbb{R}$ eingeschränkt auf $\mathbb{R} \times \mathbb{R}$ die übliche Multiplikation reeller Zahlen ist?

8. Sind die folgenden Vektoren linear unabhängig?

a) $1, \sqrt{2}, \sqrt{3}$ im $\mathbb{Q}$ -Vektorraum $\mathbb{R}$.

b) $(1, 2, 3), (4, 5, 6), (7, 8, 9)$ im $\mathbb{R}^3$.

c) $\left(\frac{1}{n+x}\right)_{n \in \mathbb{N}}$ in $\text{Abb}(\mathbb{R}_+^*, \mathbb{R})$.

d) $(\cos nx, \sin mx)_{n, m \in \mathbb{N} \setminus \{0\}}$ in $\text{Abb}(\mathbb{R}, \mathbb{R})$.

9. Für welche $t \in \mathbb{R}$ sind die folgenden Vektoren aus $\mathbb{R}^3$ linear abhängig?

$$(1, 3, 4), \quad (3, t, 11), \quad (-1, -4, 0).$$

10. Stellen Sie den Vektor w jeweils als Linearkombination der Vektoren v_1, v_2, v_3 dar:

a) $w = (6, 2, 1), v_1 = (1, 0, 1), v_2 = (7, 3, 1), v_3 = (2, 5, 8)$.

b) $w = (2, 1, 1), v_1 = (1, 5, 1), v_2 = (0, 9, 1), v_3 = (3, -3, 1)$.

1.5 Basis und Dimension

Nun muß die erste ernsthafte technische Schwierigkeit überwunden werden, um einem Vektorraum eindeutig eine Zahl (genannt Dimension) als Maß für seine Größe zuzuordnen zu können.

1.5.1. Zunächst die wichtigste

Definition. Eine Familie $\mathcal{B} = (v_i)_{i \in I}$ in einem Vektorraum V heißt *Erzeugendensystem* von V , wenn

$$V = \text{span}(v_i)_{i \in I},$$

d.h. wenn jedes $v \in V$ Linearkombination von endlich vielen v_i ist.

Eine Familie $\mathcal{B} = (v_i)_{i \in I}$ in V heißt *Basis* von V , wenn sie ein linear unabhängiges Erzeugendensystem ist.

V heißt *endlich erzeugt*, falls es ein endliches Erzeugendensystem (d.h. eine endliche Familie $\mathcal{B} = (v_1, \dots, v_n)$ mit $V = \text{span}(v_i)$ gibt. Ist $\mathcal{B}$ eine endliche Basis, so nennt man die Zahl n die *Länge der Basis*.

Beispiele. a) Die Begriffe sind so erklärt, daß die leere Familie eine Basis des Nullvektorraumes ist. Diese kleine Freude kann man ihr gönnen.

b) $\mathcal{K} := (e_1, \dots, e_n)$ ist eine Basis des K^n , sie heißt die *kanonische Basis* oder *Standardbasis* (vgl. Beispiel b) in 1.4.4).

c) Im Vektorraum $M(m \times n; K)$ hat man die Matrizen

$$E_i^j = \begin{pmatrix} & & & 0 & & & \\ & & & \vdots & & & \\ & & & 0 & & & \\ 0 & \dots & 0 & 1 & 0 & \dots & 0 \\ & & & 0 & & & \\ & & & \vdots & & & \\ & & & 0 & & & \end{pmatrix}$$

mit einer Eins in der i -ten Zeile und j -ten Spalte und sonst Nullen. Diese $m \times n$ -Matrizen bilden eine Basis von $M(m \times n; K)$. Das ist wieder nur eine Variante von Beispiel b).

d) $(1, i)$ ist eine Basis des $\mathbb{R}$ -Vektorraumes $\mathbb{C}$.

e) $(1, t, t^2, \dots)$ ist eine Basis unendlicher Länge des Polynomrings $K[t]$.

1.5.2. Das sieht alles recht einfach aus, bis auf eine zunächst spitzfindig erscheinende Frage: wenn man im K^n neben der Standardbasis irgendeine andere Basis findet, ist es gar nicht klar, daß sie die gleiche Länge hat. Im K^n wäre das noch zu verschmerzen, aber schon bei Untervektorräumen $W \subset K^n$ gibt es keine Standardbasis mehr. Es ist nicht einmal ohne weiteres klar, daß jedes solche W endlich erzeugt ist (Korollar 3 in 1.5.5). Daher kann man es nicht umgehen, die Längen verschiedener Basen zu vergleichen. Bevor wir das in Angriff nehmen, noch einige oft benutzte Varianten der Definition einer Basis. Zur Vereinfachung der Bezeichnungen betrachten wir dabei nur endliche Familien.

Satz. Für eine Familie $\mathcal{B} = (v_1, \dots, v_n)$ von Vektoren eines K -Vektorraumes $V \neq \{0\}$ sind folgende Bedingungen gleichwertig:

- i) $\mathcal{B}$ ist eine Basis, d.h. ein linear unabhängiges Erzeugendensystem.
- ii) $\mathcal{B}$ ist ein „unverkürzbares“ Erzeugendensystem, d.h.

$$(v_1, \dots, v_{r-1}, v_{r+1}, \dots, v_n)$$

ist für jedes $r \in \{1, \dots, n\}$ kein Erzeugendensystem mehr.

- iii) Zu jedem $v \in V$ gibt es eindeutig bestimmte $\lambda_1, \dots, \lambda_n \in K$ mit

$$v = \lambda_1 v_1 + \dots + \lambda_n v_n,$$

d.h. $\mathcal{B}$ ist ein Erzeugendensystem mit der zusätzlichen Eindeigkeitseigenschaft.

- iv) $\mathcal{B}$ ist „unverlängerbar“ linear unabhängig, d.h. $\mathcal{B}$ ist linear unabhängig, und für jedes $v \in V$ wird die Familie $(v_1, \dots, v_n, v)$ linear abhängig.

Beweis. i) $\Rightarrow$ ii). Gegeben sei ein Erzeugendensystem $\mathcal{B}$. Ist $\mathcal{B}$ verkürzbar, also zur Vereinfachung der Notation mit $r = 1$

$$v_1 = \lambda_2 v_2 + \dots + \lambda_n v_n, \quad \text{so folgt} \quad (-1)v_1 + \lambda_2 v_2 + \dots + \lambda_n v_n = 0.$$

Also ist $\mathcal{B}$ linear abhängig.

ii) $\Rightarrow$ iii). Sei wieder $\mathcal{B}$ ein Erzeugendensystem. Ist die Eindeigkeitseigenschaft verletzt, so gibt es ein $v \in V$ mit

$$v = \lambda_1 v_1 + \dots + \lambda_n v_n = \mu_1 v_1 + \dots + \mu_n v_n,$$

und o.B.d.A. $\lambda_1 \neq \mu_1$. Subtraktion der Linearkombinationen und Division durch $\lambda_1 - \mu_1$ ergibt

$$v_1 = \frac{\mu_2 - \lambda_2}{\lambda_1 - \mu_1} v_2 + \dots + \frac{\mu_n - \lambda_n}{\lambda_1 - \mu_1} v_n,$$

also ist $\mathcal{B}$ verkürzbar.

iii) $\Rightarrow$ iv). Aus iii) folgt, daß $\mathcal{B}$ linear unabhängig ist. (Lemma in 1.4.5). Ist $v \in V$, so ist

$$v = \lambda_1 v_1 + \dots + \lambda_n v_n, \quad \text{also} \quad \lambda_1 v_1 + \dots + \lambda_n v_n + (-1)v = 0,$$

d.h. $(v_1, \dots, v_n, v)$ ist linear abhängig.

iv) $\Rightarrow$ i). Sei $\mathcal{B}$ unverlängerbar linear unabhängig. Für jedes $v \in V$ gibt es $\lambda_1, \dots, \lambda_n, \lambda \in K$ mit

$$\lambda_1 v_1 + \dots + \lambda_n v_n + \lambda v = 0.$$

Da $\mathcal{B}$ linear unabhängig ist, muß $\lambda \neq 0$ sein, also ist

$$v = -\frac{\lambda_1}{\lambda} v_1 - \dots - \frac{\lambda_n}{\lambda} v_n,$$

und es ist bewiesen, daß $\mathcal{B}$ ein Erzeugendensystem ist. $\square$

Der Beweis von iv) $\Rightarrow$ i) ergibt den

Zusatz. Ist V nicht endlich erzeugt, so gibt es eine unendliche linear unabhängige Familie.

Beweis. Es genügt zu zeigen, daß es für beliebiges n zu linear unabhängigen Vektoren $v_1, \dots, v_n$ einen weiteren Vektor v gibt, so daß auch $(v_1, \dots, v_n, v)$ linear unabhängig ist. Wäre $(v_1, \dots, v_n, v)$ für jedes $v \in V$ linear abhängig, so wäre nach obigem Argument $(v_1, \dots, v_n)$ ein Erzeugendensystem, was der Voraussetzung widerspricht. $\square$

1.5.3. Die Bedeutung des obigen Satzes erkennt man schon an seinem gar nicht selbstverständlichen Korollar, dem

Basisauswahlsatz. Aus jedem endlichen Erzeugendensystem eines Vektorraumes kann man eine Basis auswählen. Insbesondere hat jeder endlich erzeugte Vektorraum eine endliche Basis.

Beweis. Von dem gegebenen Erzeugendensystem nehme man so lange einzelne Vektoren weg, bis es unverkürzbar geworden ist. Da am Anfang nur endlich viele da waren, führt das Verfahren zum Ziel. $\square$

Allgemeiner gilt das

Theorem. Jeder Vektorraum besitzt eine Basis.

Der Beweis ist wesentlich schwieriger, wenn es kein endliches Erzeugendensystem gibt, weil man möglicherweise unendlich viele Vektoren weglassen muß, bis die Unverkürzbarkeit erreicht ist. Ein Beweis des Theorems erfordert Hilfsmittel aus der Mengenlehre, etwa das ZORNsche Lemma. Darauf gehen wir hier nicht ein (vgl. etwa [B1], p.261).

Im Falle nicht endlich erzeugter Vektorräume sind Basen im hier definierten Sinn von geringer Bedeutung. Hier ist es für die Anwendungen in der Analysis wichtiger, konvergente unendliche Linearkombinationen zu untersuchen. Damit beschäftigt sich die Funktionalanalysis (vgl. etwa [M-V]).

1.5.4. Um die Längen verschiedener Basen zu vergleichen, muß man systematisch Vektoren austauschen. Dieses Verfahren entwickelte E. STEINITZ im Jahre 1910 bei einem ähnlichen Problem in der Theorie der Körpererweiterungen. Ein einzelner Schritt des Verfahrens wird geregelt durch das

Austauschlemma. *Gegeben sei ein K -Vektorraum V mit der Basis*

$$\mathcal{B} = (v_1, \dots, v_r) \quad \text{und} \quad w = \lambda_1 v_1 + \dots + \lambda_r v_r \in V.$$

Ist $k \in \{1, \dots, r\}$ mit $\lambda_k \neq 0$, so ist

$$\mathcal{B}' := (v_1, \dots, v_{k-1}, w, v_{k+1}, \dots, v_r)$$

wieder eine Basis von V . Man kann also v_k gegen w austauschen.

Beweis. Zur Vereinfachung der Schreibweise können wir annehmen, daß $k = 1$ ist (durch Umnummerierung kann man das erreichen). Es ist also zu zeigen, daß $\mathcal{B}' = (w, v_2, \dots, v_r)$ eine Basis von V ist. Ist $v \in V$, so ist

$$v = \mu_1 v_1 + \dots + \mu_r v_r$$

mit $\mu_1, \dots, \mu_r \in K$. Wegen $\lambda_1 \neq 0$ ist

$$v_1 = \frac{1}{\lambda_1} w - \frac{\lambda_2}{\lambda_1} v_2 - \dots - \frac{\lambda_r}{\lambda_1} v_r \quad \text{also}$$

$$v = \frac{\mu_1}{\lambda_1} w + \left(\mu_2 - \frac{\mu_1 \lambda_2}{\lambda_1} \right) v_2 + \dots + \left(\mu_r - \frac{\mu_1 \lambda_r}{\lambda_1} \right) v_r,$$

womit gezeigt ist, daß $\mathcal{B}'$ ein Erzeugendensystem ist.

Zum Nachweis der linearen Unabhängigkeit von $\mathcal{B}'$ sei

$$\mu w + \mu_2 v_2 + \dots + \mu_r v_r = 0,$$

wobei $\mu, \mu_2, \dots, \mu_r \in K$. Setzt man $w = \lambda_1 v_1 + \dots + \lambda_r v_r$ ein, so ergibt sich

$$\mu \lambda_1 v_1 + (\mu \lambda_2 + \mu_2) v_2 + \dots + (\mu \lambda_r + \mu_r) v_r = 0,$$

also $\mu \lambda_1 = \mu \lambda_2 + \mu_2 = \dots = \mu \lambda_r + \mu_r = 0$, da $\mathcal{B}$ linear unabhängig war. Wegen $\lambda_1 \neq 0$ folgt $\mu = 0$ und damit $\mu_2 = \dots = \mu_r = 0$. $\square$

Durch Iteration erhält man den

Austauschsatz. In einem K -Vektorraum V seien eine Basis

$$\mathcal{B} = (v_1, \dots, v_r)$$

und eine linear unabhängige Familie $(w_1, \dots, w_n)$ gegeben. Dann ist $n \leq r$, und es gibt Indizes $i_1, \dots, i_n \in \{1, \dots, r\}$ derart, daß man nach Austausch von v_{i_1} gegen $w_1, \dots, v_{i_n}$ gegen w_n wieder eine Basis von V erhält. Numeriert man so um, daß $i_1 = 1, \dots, i_n = n$ ist, so bedeutet das, daß

$$\mathcal{B}^* = (w_1, \dots, w_n, v_{n+1}, \dots, v_r)$$

eine Basis von V ist.

Vorsicht! Die Ungleichung $n \leq r$ wird nicht vorausgesetzt, sondern gefolgert.

Beweis durch Induktion nach n . Für $n = 0$ ist nichts zu beweisen. Sei also $n \geq 1$, und sei der Satz schon für $n - 1$ bewiesen (Induktionsannahme).

Da auch $(w_1, \dots, w_{n-1})$ linear unabhängig ist, ergibt die Induktionsannahme, daß (bei geeigneter Numerierung) $(w_1, \dots, w_{n-1}, v_n, \dots, v_r)$ eine Basis von V ist. Da nach Induktionsannahme $n - 1 \leq r$ gilt, muß zum Nachweis von $n \leq r$ nur noch der Fall $n - 1 = r$ ausgeschlossen werden. Dann wäre aber $(w_1, \dots, w_{n-1})$ schon eine Basis von V , was Aussage iv) aus Satz 1.5.2 widerspricht. Wir schreiben

$$w_n = \lambda_1 w_1 + \dots + \lambda_{n-1} w_{n-1} + \lambda_n v_n + \dots + \lambda_r v_r$$

mit $\lambda_1, \dots, \lambda_r \in K$. Wäre $\lambda_n = \dots = \lambda_r = 0$, so hätte man einen Widerspruch zur linearen Unabhängigkeit von $w_1, \dots, w_n$. Bei erneuter geeigneter Numerierung können wir also $\lambda_n \neq 0$ annehmen, und wie wir im Austauschlemma gesehen haben, läßt sich daher v_n gegen w_n austauschen. Also ist $\mathcal{B}^*$ eine Basis von V . $\square$

1.5.5. Nach Überwindung dieser kleinen technischen Schwierigkeiten läuft die Theorie wieder wie von selbst. Wir notieren die wichtigsten Folgerungen.

Korollar 1. Hat ein K -Vektorraum V eine endliche Basis, so ist jede Basis von V endlich.

Beweis. Sei $(v_1, \dots, v_r)$ eine endliche Basis und $(w_i)_{i \in I}$ eine beliebige Basis von V . Wäre I unendlich, so gäbe es $i_1, \dots, i_{r+1} \in I$ derart, daß $w_{i_1}, \dots, w_{i_{r+1}}$ linear unabhängig wären. Das widerspricht aber dem Austauschsatz. $\square$

Korollar 2. Je zwei endliche Basen eines K -Vektorraumes haben gleiche Länge.

Beweis. Sind $(v_1, \dots, v_r)$ und $(w_1, \dots, w_k)$ zwei Basen, so kann man den Austauschsatz zweimal anwenden, was $k \leq r$ und $r \leq k$, also $r = k$ ergibt. $\square$

Mit Hilfe dieser Ergebnisse können wir nun in sinnvoller Weise die Dimension eines Vektorraumes erklären.

Definition. Ist V ein K -Vektorraum, so definieren wir

$$\dim_K V := \begin{cases} \infty, & \text{falls } V \text{ keine endliche Basis besitzt,} \\ r, & \text{falls } V \text{ eine Basis der Länge } r \text{ besitzt.} \end{cases}$$

$\dim_K V$ heißt die *Dimension* von V über K . Falls klar ist, welcher Körper gemeint ist, schreibt man auch $\dim V$.

Korollar 3. Ist $W \subset V$ Untervektorraum eines endlich erzeugten Vektorraumes V , so ist auch W endlich erzeugt, und es gilt $\dim W \leq \dim V$.

Aus $\dim W = \dim V$ folgt $W = V$.

Beweis. Wäre W nicht endlich erzeugt, so gäbe es nach dem Zusatz aus 1.5.2 eine unendliche linear unabhängige Familie, was dem Austauschsatz widerspricht. Also hat W eine endliche Basis, und wieder nach dem Austauschsatz ist ihre Länge höchstens gleich $\dim V$.

Sei $n = \dim W = \dim V$ und $w_1, \dots, w_n$ Basis von W . Ist $W \neq V$, so gibt es ein $v \in V \setminus W$ und $w_1, \dots, w_n, v$ sind linear unabhängig im Widerspruch zum Austauschsatz. $\square$

In 1.5.3 hatten wir gesehen, daß man aus einem endlichen Erzeugendensystem eine Basis auswählen kann. Manchmal ist die Konstruktion „aus der anderen Richtung“ wichtig:

Basisergänzungssatz. In einem endlich erzeugten Vektorraum V seien linear unabhängige Vektoren $w_1, \dots, w_n$ gegeben. Dann kann man $w_{n+1}, \dots, w_r$ finden, so daß

$$\mathcal{B} = (w_1, \dots, w_n, w_{n+1}, \dots, w_r)$$

eine Basis von V ist.

Beweis. Sei $(v_1, \dots, v_m)$ ein Erzeugendensystem. Nach 1.5.3 kann man daraus eine Basis auswählen, etwa $(v_1, \dots, v_r)$ mit $r \leq m$. Nun wendet man den Austauschsatz an und sieht, daß bei geeigneter Numerierung durch

$$w_{n+1} := v_{n+1}, \dots, w_r := v_r$$

die gesuchte Ergänzung gefunden ist. $\square$

Beispiele. a) $\dim K^n = n$, denn K^n hat die kanonische Basis $(e_1, \dots, e_n)$. Nach Korollar 2 hat auch jede andere Basis von K^n die Länge n , was gar nicht selbstverständlich ist.

b) Geraden (bzw. Ebenen) durch den Nullpunkt des $\mathbb{R}^n$ sind Untervektorräume der Dimension 1 (bzw. 2).

- c) Für den Polynomring gilt $\dim_K K[t] = \infty$.
 d) $\dim_{\mathbb{R}} \mathbb{C} = 2$ denn 1 und i bilden eine Basis. Dagegen ist $\dim_{\mathbb{C}} \mathbb{C} = 1$.
 e) $\dim_{\mathbb{Q}} \mathbb{R} = \infty$ (vgl. Aufgabe 4).

1.5.6. Bei der Definition eines Vektorraumes in 1.4.1 hatten wir einen Körper K zugrundegelegt. Zur Formulierung der Axiome genügt ein kommutativer Ring R mit Einselement, man spricht dann von einem *Modul über R* . Die Begriffe wie Linearkombination, Erzeugendensystem und lineare Unabhängigkeit kann man in dieser allgemeineren Situation analog erklären.

In den vorangegangenen Beweisen wird immer wieder durch Skalare dividiert, was einen Skalarenkörper voraussetzt. Über einem Ring ist von den erhaltenen Aussagen über Basis und Dimension wenig zu retten. Daher beschränken wir uns auf zwei Aufgaben (8 und 9), die als Warnung vor diesen Gefahren dienen sollen.

1.5.7. Im Basisauswahlsatz 1.5.3 hatten wir bewiesen, daß man aus jedem endlichen Erzeugendensystem eine Basis auswählen kann. Für die Praxis ist das Verfahren des Weglassens (und die Kontrolle, ob ein Erzeugendensystem übrig bleibt) nicht gut geeignet. Weitaus einfacher ist es, aus einem Erzeugendensystem eine Basis linear zu kombinieren. Wir behandeln hier den Spezialfall eines Untervektorraumes $W \subset K^n$; in 2.4.2 werden wir sehen, daß sich der allgemeine Fall darauf zurückführen läßt.

Seien also $a_1, \dots, a_m \in K^n$ gegeben, und sei $W = \text{span}(a_1, \dots, a_m)$. Sind die Vektoren a_i Zeilen, so ergeben sie untereinandergeschrieben eine Matrix

$$A = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix} \in M(m \times n; K),$$

d.h. es ist $a_i = (a_{i1}, \dots, a_{in})$.

Beispiel. Aus der kanonischen Basis $(e_1, \dots, e_n)$ von K^n erhält man die Matrix

$$E_n := \begin{pmatrix} 1 & & 0 \\ & \ddots & \\ 0 & & 1 \end{pmatrix} \in M(n \times n; K),$$

man nennt sie die n -reihige *Einheitsmatrix*. Dieser Name ist durch ihre Wirkung bei der Matrizenmultiplikation erklärt (vgl. 2.5.4). Die Einträge von

$E_n = (\delta_{ij})$ sind die sogenannten **KRONECKER-Symbole**

$$\delta_{ij} := \begin{cases} 0 & \text{für } i \neq j, \\ 1 & \text{für } i = j. \end{cases}$$

Nun kommen wir zurück auf die schon in Kapitel 0 benutzten Zeilenumformungen. Anstatt der reellen Zahlen stehen Einträge aus einem beliebigen Körper K , und wir betrachten vier verschiedene Arten von *elementaren Zeilenumformungen*:

I Multiplikation der i -ten Zeile mit $\lambda \in K^*$:

$$A = \begin{pmatrix} \vdots \\ a_i \\ \vdots \end{pmatrix} \mapsto \begin{pmatrix} \vdots \\ \lambda a_i \\ \vdots \end{pmatrix} =: A_I.$$

II Addition der j -ten Zeile zur i -ten Zeile:

$$A = \begin{pmatrix} \vdots \\ a_i \\ \vdots \\ a_j \\ \vdots \end{pmatrix} \mapsto \begin{pmatrix} \vdots \\ a_i + a_j \\ \vdots \\ a_j \\ \vdots \end{pmatrix} =: A_{II}.$$

III Addition der λ -fachen j -ten Zeile zur i -ten Zeile ($\lambda \in K^*$):

$$A = \begin{pmatrix} \vdots \\ a_i \\ \vdots \\ a_j \\ \vdots \end{pmatrix} \mapsto \begin{pmatrix} \vdots \\ a_i + \lambda a_j \\ \vdots \\ a_j \\ \vdots \end{pmatrix} =: A_{III}.$$

IV Vertauschen der i -ten Zeile mit der j -ten Zeile:

$$A = \begin{pmatrix} \vdots \\ a_i \\ \vdots \\ a_j \\ \vdots \end{pmatrix} \mapsto \begin{pmatrix} \vdots \\ a_j \\ \vdots \\ a_i \\ \vdots \end{pmatrix} =: A_{IV}.$$

Dabei bezeichnen jeweils $a_1, \dots, a_m$ die Zeilen von A , es ist stets $i \neq j$ vorausgesetzt, und an den mit Punkten markierten Zeilen ändert sich nichts.

Die Typen III und IV entsprechen 1) und 2) aus 0.4.6. Die Typen I und II sind noch elementarer, denn man kann III und IV daraus durch Kombination erhalten,

und zwar nach folgendem Rezept:

$$\begin{pmatrix} a_i \\ a_j \end{pmatrix} \xrightarrow{\text{I}} \begin{pmatrix} a_i \\ \lambda a_j \end{pmatrix} \xrightarrow{\text{II}} \begin{pmatrix} a_i + \lambda a_j \\ \lambda a_j \end{pmatrix} \xrightarrow{\text{I}} \begin{pmatrix} a_i + \lambda a_j \\ a_j \end{pmatrix} \quad \text{bzw.}$$

$$\begin{pmatrix} a_i \\ a_j \end{pmatrix} \xrightarrow{\text{I}} \begin{pmatrix} a_i \\ -a_j \end{pmatrix} \xrightarrow{\text{II}} \begin{pmatrix} a_i \\ a_i - a_j \end{pmatrix} \xrightarrow{\text{III}}$$

$$\begin{pmatrix} a_i - (a_i - a_j) \\ a_i - a_j \end{pmatrix} = \begin{pmatrix} a_j \\ a_i - a_j \end{pmatrix} \xrightarrow{\text{II}} \begin{pmatrix} a_j \\ a_i \end{pmatrix}.$$

Zum Verständnis der Wirkung von Zeilenumformungen hilft ein weiterer Begriff:

Definition. Ist $A \in M(m \times n; K)$ mit Zeilen $a_1, \dots, a_m$, so heißt

$$\text{ZR}(A) := \text{span}(a_1, \dots, a_m) \subset K^n$$

der Zeilenraum von A .

Lemma. Ist B aus A durch elementare Zeilenumformungen entstanden, so ist $\text{ZR}(B) = \text{ZR}(A)$.

Beweis. Nach der obigen Bemerkung genügt es, die Typen I und II zu betrachten. Ist $B = A_{\text{I}}$ und $v \in \text{ZR}(A)$, so ist

$$v = \dots + \mu_i a_i + \dots = \dots + \frac{\mu_i}{\lambda} (\lambda a_i) + \dots,$$

also auch $v \in \text{ZR}(B)$. Analog folgt $v \in \text{ZR}(A)$ aus $v \in \text{ZR}(B)$.

Ist $B = A_{\text{II}}$ und $v \in \text{ZR}(A)$, so ist

$$v = \dots + \mu_i a_i + \dots + \mu_j a_j + \dots = \dots + \mu_i (a_i + a_j) + \dots + (\mu_j - \mu_i) a_j + \dots,$$

also $v \in \text{ZR}(B)$ und analog umgekehrt. $\square$

Wie in 0.4.7 beweist man den

Satz. Jede Matrix $A \in M(m \times n; K)$ kann man durch elementare Zeilenumformungen auf Zeilenstufenform bringen. $\square$

Damit ist das zu Beginn dieses Abschnitts formulierte Problem gelöst: Hat man aus den gegebenen Vektoren $a_1, \dots, a_m$ die Matrix A aufgestellt und diese zu B in Zeilenstufenform umgeformt, so sind die von Null verschiedenen Zeilen $b_1, \dots, b_r$ von B eine Basis von $W = \text{ZR}(A) = \text{ZR}(B)$, denn $b_1, \dots, b_r$ sind nach Beispiel b) in 1.4.5 linear unabhängig.

Beispiel. Im $\mathbb{R}^5$ seien die Vektoren

$$\begin{aligned}a_1 &= (0, 0, 0, 2, -1), \\a_2 &= (0, 1, -2, 1, 0), \\a_3 &= (0, -1, 2, 1, -1), \\a_4 &= (0, 0, 0, 1, 2)\end{aligned}$$

gegeben. Dann verläuft die Rechnung wie folgt:

$$\begin{aligned}A &= \begin{array}{ccccc|ccccc} 0 & 0 & 0 & 2 & -1 & 0 & 1 & -2 & 1 & 0 \\ 0 & 1 & -2 & 1 & 0 & 0 & 0 & 0 & 2 & -1 \\ 0 & -1 & 2 & 1 & -1 & 0 & -1 & 2 & 1 & -1 \\ 0 & 0 & 0 & 1 & 2 & 0 & 0 & 0 & 1 & 2 \end{array} \rightsquigarrow \begin{array}{ccccc|ccccc} 0 & 1 & -2 & 1 & 0 & 0 & 0 & 0 & 2 & -1 \\ 0 & 0 & 0 & 1 & 2 & 0 & -1 & 2 & 1 & -1 \\ 0 & 0 & 0 & 1 & 2 & 0 & 0 & 0 & 1 & 2 \end{array} \\ &\rightsquigarrow \begin{array}{ccccc|ccccc} 0 & 1 & -2 & 1 & 0 & 0 & 0 & 0 & 1 & 2 \\ 0 & 0 & 0 & 2 & -1 & 0 & 0 & 0 & 2 & -1 \\ 0 & 0 & 0 & 2 & -1 & 0 & 0 & 0 & 2 & -1 \\ 0 & 0 & 0 & 1 & 2 & 0 & 0 & 0 & 2 & -1 \end{array} \\ &\rightsquigarrow \begin{array}{ccccc|ccccc} 0 & 1 & -2 & 1 & 0 & 0 & 1 & -2 & 1 & 0 \\ 0 & 0 & 0 & 1 & 2 & 0 & 0 & 0 & 1 & 2 \\ 0 & 0 & 0 & 0 & -5 & 0 & 0 & 0 & 0 & -5 \\ 0 & 0 & 0 & 0 & -5 & 0 & 0 & 0 & 0 & 0 \end{array} = B.\end{aligned}$$

Also ist eine Basis von $W = \text{span}(a_1, a_2, a_3, a_4)$ gegeben durch

$$\begin{aligned}b_1 &= (0, 1, -2, 1, 0), \\b_2 &= (0, 0, 0, 1, 2), \\b_3 &= (0, 0, 0, 0, -5).\end{aligned}$$

1.5.8. Ob man Vektoren als Zeilen oder Spalten schreibt ist willkürlich, aber der Übergang von der einen zur anderen Konvention wirft Fragen auf. Macht man in einer Matrix Zeilen zu Spalten, so werden Spalten zu Zeilen; man nennt das *Transposition*: Zu

$$A = (a_{ij}) \in M(m \times n; K) \quad \text{ist} \quad {}^t A = (a'_{ij}) \in M(n \times m; K) \quad \text{mit} \quad a'_{ij} := a_{ji}.$$

$'A$ heißt die *Transponierte* von A . Zum Beispiel ist für $m = 2$ und $n = 3$

$$\begin{pmatrix} 1 & 0 & 3 \\ 0 & 2 & 1 \end{pmatrix}' = \begin{pmatrix} 1 & 0 \\ 0 & 2 \\ 3 & 1 \end{pmatrix}.$$

Eine quadratische Matrix wird dabei an der Diagonale gespiegelt. Abstrakt gesehen ist die Transposition eine bijektive Abbildung

$$M(m \times n; K) \rightarrow M(n \times m; K), \quad A \mapsto 'A,$$

und es gelten die folgenden Rechenregeln:

- 1) $'(A + B) = 'A + 'B$,
- 2) $'(\lambda \cdot A) = \lambda \cdot 'A$,
- 3) $'('A) = A$.

Nun kann man ganz analog zu 1.5.7 für eine Matrix A elementare Spaltenumformungen, Spaltenraum $\text{SR}(A)$, Spaltenstufenform etc. erklären, was durch Transposition auf die entsprechenden Begriffe für Zeilen zurückgeführt werden kann. Das sind einfache Spielereien, aber ein ernsthaftes Problem wird deutlich durch die folgende

Definition. Für eine Matrix $A \in M(m \times n; K)$ sei

$$\begin{aligned} \text{Zeilenrang } A &:= \dim \text{ZR}(A) \quad \text{und} \\ \text{Spaltenrang } A &:= \dim \text{SR}(A). \end{aligned}$$

Man beachte dabei, daß Zeilen- und Spaltenraum in verschiedenen Vektorräumen liegen:

$$\text{ZR}(A) \subset K^n \quad \text{und} \quad \text{SR}(A) \subset K^m.$$

Zum Beispiel für $A = E_n$ ist $\text{ZR}(E_n) = \text{SR}(E_n) = K^n$, also

$$\text{Zeilenrang } E_n = n = \text{Spaltenrang } E_n.$$

Bei der Behandlung linearer Gleichungssysteme benötigt man die etwas überraschende Tatsache, daß diese beiden Zahlen für jede Matrix gleich sind. In 2.6.6 steht genügend viel Theorie für einen indexfreien Beweis zur Verfügung, in Kapitel 6 wird der abstrakte Hintergrund beleuchtet. Der Schlüssel für einen direkten Beweis mit den Hilfsmitteln dieses Kapitels ist das

Lemma. In der Matrix $A \in M(m \times n; K)$ sei die letzte Zeile Linearkombination der vorhergehenden, $\bar{A} \in M(m-1 \times n; K)$ entstehe aus A durch weglassen der letzten Zeile. Dann ist

$$\text{Spaltenrang } \bar{A} = \text{Spaltenrang } A.$$

Beweis. Wir bezeichnen die Zeilen von A mit $a_1, \dots, a_m$ und die Spalten mit $a^1, \dots, a^n$ (man beachte, daß der obere Index keine Potenz bedeuten soll). Nach Voraussetzung gibt es $\mu_1, \dots, \mu_{m-1} \in K$ mit

$$a_m = \mu_1 a_1 + \dots + \mu_{m-1} a_{m-1}.$$

Das bedeutet, daß alle Spalten von A enthalten sind im Untervektorraum

$$W = \{(x_1, \dots, x_m) \in K^m : x_m = \mu_1 x_1 + \dots + \mu_{m-1} x_{m-1}\},$$

also ist $\text{SR}(A) \subset W$. Nun betrachten wir die „Projektion“

$$\pi : W \rightarrow K^{m-1}, \quad x = (x_1, \dots, x_m) \mapsto \bar{x} = (x_1, \dots, x_{m-1}).$$

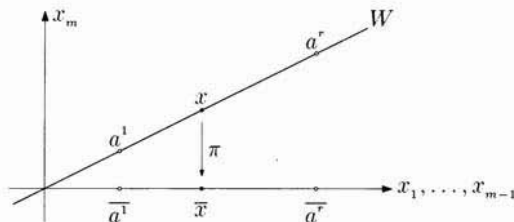


Bild 1.7

Nach Definition von $\bar{A}$ gilt $\text{SR}(\bar{A}) = \pi(\text{SR}(A))$, zu zeigen ist $\dim \text{SR}(\bar{A}) = \dim \text{SR}(A)$.

Nach 1.5.3 können wir annehmen, daß $(a^1, \dots, a^r)$ für ein r mit $0 \leq r \leq n$ eine Basis von $\text{SR}(A)$ ist. Offensichtlich ist $\bar{a}^1, \dots, \bar{a}^r$ ein Erzeugendensystem von $\text{SR}(\bar{A})$. Ist

$$\lambda_1 \bar{a}^1 + \dots + \lambda_r \bar{a}^r = 0, \quad (*)$$

so folgt daraus für die m -ten Komponenten von $a^1, \dots, a^r$

$$\begin{aligned} \lambda_1 a_m^1 + \dots + \lambda_r a_m^r &= \lambda_1 \sum_{i=1}^{m-1} \mu_i a_i^1 + \dots + \lambda_r \sum_{i=1}^{m-1} \mu_i a_i^r \\ &= \mu_1 \sum_{j=1}^r \lambda_j a_j^1 + \dots + \mu_{m-1} \sum_{j=1}^r \lambda_j a_j^{m-1} \\ &= \mu_1 \cdot 0 + \dots + \mu_{m-1} \cdot 0 = 0, \end{aligned}$$

also folgt aus (*), daß $\lambda_1 a^1 + \dots + \lambda_r a^r = 0$ und somit $\lambda_1 = \dots = \lambda_r = 0$. Daher ist $\bar{a}^1, \dots, \bar{a}^r$ eine Basis von $\text{SR}(\bar{A})$.

Selbstverständlich bleibt der Zeilenraum ganz unverändert. Das folgt etwa aus dem Lemma in 1.5.7. $\square$

Der Rest des Beweises für die Gleichheit von Zeilenrang und Spaltenrang ist Aufgabe 13.

Aufgaben zu 1.5

1. Gegeben seien im $\mathbb{R}^5$ die Vektoren $v_1 = (4, 1, 1, 0, -2)$, $v_2 = (0, 1, 4, -1, 2)$, $v_3 = (4, 3, 9, -2, 2)$, $v_4 = (1, 1, 1, 1, 1)$, $v_5 = (0, -2, -8, 2, -4)$.

- Bestimmen Sie eine Basis von $V = \text{span}(v_1, \dots, v_5)$.
- Wählen Sie alle möglichen Basen von V aus den Vektoren $v_1, \dots, v_5$ aus, und kombinieren Sie jeweils $v_1, \dots, v_5$ daraus linear.

2. Geben Sie für folgende Vektorräume jeweils eine Basis an:

- $\{(x_1, x_2, x_3) \in \mathbb{R}^3 : x_1 = x_3\}$,
- $\{(x_1, x_2, x_3, x_4) \in \mathbb{R}^4 : x_1 + 3x_2 + 2x_4 = 0, 2x_1 + x_2 + x_3 = 0\}$,
- $\text{span}(t^2, t^2 + t, t^2 + 1, t^2 + t + 1, t^7 + t^5) \subset \mathbb{R}[t]$,
- $\{f \in \text{Abb}(\mathbb{R}, \mathbb{R}) : f(x) = 0 \text{ bis auf endlich viele } x \in \mathbb{R}\}$.

3. Für $d \in \mathbb{N}$ sei

$K[t_1, \dots, t_n]_{(d)} := \{F \in K[t_1, \dots, t_n] : F \text{ ist homogen vom Grad } d \text{ oder } F = 0\}$
(vgl. Aufgabe 9 zu 1.3). Beweisen Sie, dass $K[t_1, \dots, t_n]_{(d)} \subset K[t_1, \dots, t_n]$ ein Untervektorraum ist und bestimmen Sie $\dim K[t_1, \dots, t_n]_{(d)}$.

4. Zeigen Sie, dass $\mathbb{C}$ endlich erzeugt über $\mathbb{R}$ ist, aber $\mathbb{R}$ nicht endlich erzeugt über $\mathbb{Q}$.

5. Ist $(v_i)_{i \in I}$ eine Basis des Vektorraumes V und $(w_j)_{j \in J}$ eine Basis des Vektorraumes W , so ist $((v_i, 0))_{i \in I} \cup ((0, w_j))_{j \in J}$ eine Basis von $V \times W$ (vgl. Aufgabe 2 zu 1.4). Insbesondere gilt

$$\dim V \times W = \dim V + \dim W,$$

falls $\dim V, \dim W < \infty$.

6. Sei V ein reeller Vektorraum und $a, b, c, d, e \in V$. Zeigen Sie, dass die folgenden Vektoren linear abhängig sind:

$$\begin{aligned} v_1 &= a + b + c, & v_2 &= 2a + 2b + 2c - d, & v_3 &= a - b - e, \\ v_4 &= 5a + 6b - c + d + e, & v_5 &= a - c + 3e, & v_6 &= a + b + d + e. \end{aligned}$$

7. Für einen endlichdimensionalen Vektorraum V definieren wir

$$h(V) := \sup \{n \in \mathbb{N} : \text{es gibt eine Kette } V_0 \subset V_1 \subset \dots \subset V_{n-1} \subset V_n \\ \text{von Untervektorräumen } V_i \subset V\}.$$

Zeigen Sie $h(V) = \dim V$.

8. Sei $R = \mathcal{C}(\mathbb{R}, \mathbb{R})$ der Ring der stetigen Funktionen und

$$W := \{f \in R : \text{es gibt ein } \varrho \in \mathbb{R} \text{ mit } f(x) = 0 \text{ für } x \geq \varrho\} \subset R.$$

Für $k \in \mathbb{N}$ definieren wir die Funktion

$$f_k(x) := \begin{cases} 0 & \text{für alle } x \geq k, \\ k - x & \text{für } x \leq k. \end{cases}$$

a) $W = \text{span}_R(f_k)_{k \in \mathbb{N}}$.

b) W ist über R nicht endlich erzeugt (aber R ist über R endlich erzeugt).

c) Ist die Familie $(f_k)_{k \in \mathbb{N}}$ linear abhängig über R ?

9. Zeigen Sie $\mathbb{Z} = 2\mathbb{Z} + 3\mathbb{Z}$ und folgern Sie daraus, dass es in $\mathbb{Z}$ unverkürzbare Erzeugendensysteme verschiedener Längen gibt.

10. Wie viele Elemente hat ein endlichdimensionaler Vektorraum über einem endlichen Körper?

11.* a) Ist K ein Körper mit $\text{char } K = p > 0$, so enthält K einen zu $\mathbb{Z}/p\mathbb{Z}$ isomorphen Körper und kann somit als $\mathbb{Z}/p\mathbb{Z}$ -Vektorraum aufgefasst werden.

b) Zeigen Sie: Ist K ein endlicher Körper mit $\text{char } K = p$, so hat K genau p^n Elemente, wobei $n = \dim_{\mathbb{Z}/p\mathbb{Z}} K$.

12. Zeigen Sie: Zeilenrang = Spaltenrang für Matrizen mit sehr kleiner Zeilenzahl (etwa $m = 1, 2$) und beliebig großer Spaltenzahl n .

13. Folgern Sie aus Lemma 1.5.8, dass für eine Matrix $A \in M(m \times n; K)$

a) Zeilenrang $A \leq$ Spaltenrang A ,

b) Zeilenrang $A \geq$ Spaltenrang A ,

und somit insgesamt Zeilenrang $A =$ Spaltenrang A gilt.

1.6 Summen von Vektorräumen*

Für die Menge der Linearkombinationen aus Vektoren $v_1, \dots, v_r$ hatten wir in 1.4.4 auch die Schreibweise

$$Kv_1 + \dots + Kv_r$$

angegeben. Das soll andeuten, daß dies die Menge der Summen von Vektoren aus den für $v_j \neq 0$ eindimensionalen Räumen Kv_j ist. Wir betrachten nun den Fall, daß die Summanden aus beliebigen Untervektorräumen stammen, das wird in Kapitel 4 nützlich sein. Einem eiligen Leser raten wir, dies zunächst zu überblättern.

1.6.1. Ausgangspunkt ist die folgende

Definition. Gegeben sei ein K -Vektorraum V mit Untervektorräumen $W_1, \dots, W_r \subset V$. Dann heißt

$$W_1 + \dots + W_r := \{v \in V : \text{es gibt } v_j \in W_j \text{ mit } v = v_1 + \dots + v_r\}$$

die *Summe* von $W_1, \dots, W_r$.

Ohne Schwierigkeit beweist man die

Bemerkung. Für die oben definierte Summe gilt:

a) $W_1 + \dots + W_r \subset V$ ist ein Untervektorraum.

b) $W_1 + \dots + W_r = \text{span}(W_1 \cup \dots \cup W_r)$.

c) $\dim(W_1 + \dots + W_r) \leq \dim W_1 + \dots + \dim W_r$. □

Nun ist die Frage naheliegend, wie unscharf die letzte Ungleichung ist. Im Fall $r = 2$ ist das einfach zu beantworten:

Dimensionsformel für Summen. Für endlichdimensionale Untervektorräume $W_1, W_2 \subset V$ gilt

$$\dim(W_1 + W_2) = \dim W_1 + \dim W_2 - \dim(W_1 \cap W_2).$$

Beweis. Wir beginnen mit einer Basis $(v_1, \dots, v_m)$ von $W_1 \cap W_2$ und ergänzen sie entsprechend 1.5.5 zu Basen

$$(v_1, \dots, v_m, w_1, \dots, w_k) \text{ von } W_1 \text{ und } (v_1, \dots, v_m, w'_1, \dots, w'_l) \text{ von } W_2.$$

Die Behauptung ist bewiesen, wenn wir gezeigt haben, daß

$$\mathcal{B} := (v_1, \dots, v_m, w_1, \dots, w_k, w'_1, \dots, w'_l)$$

eine Basis von $W_1 + W_2$ ist. Daß $W_1 + W_2$ von $\mathcal{B}$ erzeugt wird ist klar. Zum Beweis der linearen Unabhängigkeit sei

$$\lambda_1 v_1 + \dots + \lambda_m v_m + \mu_1 w_1 + \dots + \mu_k w_k + \mu'_1 w'_1 + \dots + \mu'_l w'_l = 0. \quad (*)$$

Setzen wir

$$v := \lambda_1 v_1 + \dots + \lambda_m v_m + \mu_1 w_1 + \dots + \mu_k w_k,$$

so ist $v \in W_1$ und $-v = \mu'_1 w'_1 + \dots + \mu'_l w'_l \in W_2$, also $v \in W_1 \cap W_2$. Also ist

$$v = \lambda'_1 v_1 + \dots + \lambda'_m v_m$$

mit $\lambda'_1, \dots, \lambda'_m \in K$, und wegen der Eindeutigkeit der Linearkombinationen folgt insbesondere $\mu_1 = \dots = \mu_k = 0$. Setzt man das in (*) ein, so folgt auch

$$\lambda_1 = \dots = \lambda_m = \mu'_1 = \dots = \mu'_l = 0. \quad \square$$

1.6.2. Der Korrekturterm in der Dimensionsformel wird durch die Dimension des Durchschnitts verursacht. Diesen Mangel einer Summendarstellung kann man auch anders charakterisieren.

Lemma. Ist $V = W_1 + W_2$, so sind folgende Bedingungen äquivalent:

- i) $W_1 \cap W_2 = \{0\}$.
- ii) Jedes $v \in V$ ist eindeutig darstellbar als $v = w_1 + w_2$ mit $w_1 \in W_1$, $w_2 \in W_2$.
- iii) Zwei von Null verschiedene Vektoren $w_1 \in W_1$ und $w_2 \in W_2$ sind linear unabhängig.

Beweis. i) $\Rightarrow$ ii): Ist $v = w_1 + w_2 = \tilde{w}_1 + \tilde{w}_2$, so folgt

$$w_1 - \tilde{w}_1 = \tilde{w}_2 - w_2 \in W_1 \cap W_2.$$

ii) $\Rightarrow$ iii): Sind w_1, w_2 linear abhängig, so erhält man verschiedene Darstellungen des Nullvektors.

iii) $\Rightarrow$ i): Ist $0 \neq v \in W_1 \cap W_2$, so erhält man einen Widerspruch zu iii) durch

$$1v + (-1)v = 0. \quad \square$$

Da Bedingung i) am kürzesten aufzuschreiben ist, die folgende

Definition. Ein Vektorraum V heißt *direkte Summe* von zwei Untervektorräumen W_1 und W_2 , in Zeichen

$$V = W_1 \oplus W_2, \quad \text{wenn} \quad V = W_1 + W_2 \quad \text{und} \quad W_1 \cap W_2 = \{0\}.$$

1.6.3. Im endlichdimensionalen Fall hat man einfacher nachzuprüfende Bedingungen für die Direktheit einer Summe.

Satz. Ist V endlichdimensional mit Untervektorräumen W_1 und W_2 , so sind folgende Bedingungen gleichwertig:

- i) $V = W_1 \oplus W_2$.
- ii) Es gibt Basen $(w_1, \dots, w_k)$ von W_1 und $(w'_1, \dots, w'_l)$ von W_2 , so daß $(w_1, \dots, w_k, w'_1, \dots, w'_l)$ eine Basis von V ist.
- iii) $V = W_1 + W_2$ und $\dim V = \dim W_1 + \dim W_2$.

Beweis. i) $\Rightarrow$ ii) $\Rightarrow$ iii) folgt aus der Dimensionsformel (einschließlich Beweis) in 1.6.1 im Spezialfall $W_1 \cap W_2 = \{0\}$.

iii) $\Rightarrow$ i): Nach der Dimensionsformel ist $\dim(W_1 \cap W_2) = 0$, also

$$W_1 \cap W_2 = \{0\}. \quad \square$$

Daraus folgt sofort die Existenz von „direkten Summanden“:

Korollar. Ist V endlichdimensional und $W \subset V$ Untervektorraum, so gibt es dazu einen (im allgemeinen nicht eindeutig bestimmten) Untervektorraum $W' \subset V$, so daß

$$V = W \oplus W'.$$

W' heißt *direkter Summand* von V zu W .

Beweis. Man nehme eine Basis $(v_1, \dots, v_m)$ von W , ergänze sie nach 1.5.4 zu einer Basis $(v_1, \dots, v_m, v_{m+1}, \dots, v_n)$ von V und definiere

$$W' := \text{span}(v_{m+1}, \dots, v_n). \quad \square$$

1.6.4. In Kapitel 4 werden wir direkte Summen von mehreren Unterräumen antreffen. Zur Vorsorge dafür die

Definition. Ein Vektorraum V heißt *direkte Summe* von Untervektorräumen $W_1, \dots, W_k$, in Zeichen

$$V = W_1 \oplus \dots \oplus W_k,$$

wenn folgende Bedingungen erfüllt sind:

DS1 $V = W_1 + \dots + W_k$.

DS2 Sind $w_1 \in W_1, \dots, w_k \in W_k$ gegeben mit $w_1 + \dots + w_k = 0$, so folgt $w_1 = \dots = w_k = 0$.

Vorsicht! Bedingung DS2 darf man für $k > 2$ nicht ersetzen durch

$$W_1 \cap \dots \cap W_k = \{0\} \quad \text{oder} \quad W_i \cap W_j = \{0\} \quad \text{für alle } i \neq j$$

(vgl. Aufgabe 1).

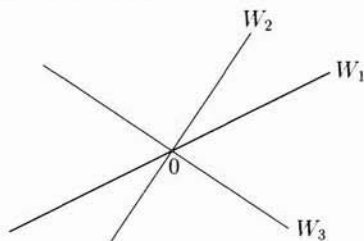


Bild 1.8

Beispiel. Ist $(v_1, \dots, v_n)$ eine Basis von V , so ist $V = K v_1 \oplus \dots \oplus K v_n$.

Satz. Für Untervektorräume $W_1, \dots, W_k$ eines endlichdimensionalen Vektorraumes V sind folgende Bedingungen äquivalent:

- i) $V = W_1 \oplus \dots \oplus W_k$.
- ii) Ist für jedes $i \in \{1, \dots, k\}$ eine Basis $(v_1^{(i)}, \dots, v_{r_i}^{(i)})$ von W_i gegeben, so ist

$$\mathcal{B} := (v_1^{(1)}, \dots, v_{r_1}^{(1)}, \dots, v_1^{(k)}, \dots, v_{r_k}^{(k)})$$

eine Basis von V .

- iii) $V = W_1 + \dots + W_k$ und $\dim V = \dim W_1 + \dots + \dim W_k$.

Man beachte die Klammern bei den oberen Indizes. Sie dienen zur Unterscheidung von Exponenten, deren Stammpfad an dieser Stelle ist.

Beweis. i) $\Rightarrow$ ii). Offensichtlich ist $\mathcal{B}$ ein Erzeugendensystem. Zum Beweis der linearen Unabhängigkeit sei

$$\mu_1^{(1)} v_1^{(1)} + \dots + \mu_{r_1}^{(1)} v_{r_1}^{(1)} + \dots + \mu_1^{(k)} v_1^{(k)} + \dots + \mu_{r_k}^{(k)} v_{r_k}^{(k)} = 0.$$

Setzen wir $w_i := \mu_1^{(i)} v_1^{(i)} + \dots + \mu_{r_i}^{(i)} v_{r_i}^{(i)}$, so bedeutet das

$$w_1 + \dots + w_k = 0,$$

und wegen DS2 muß $w_1 = \dots = w_k = 0$ sein. Also ist

$$\mu_1^{(i)} v_1^{(i)} + \dots + \mu_{r_i}^{(i)} v_{r_i}^{(i)} = 0 \quad \text{für } i = 1, \dots, k,$$

und daraus folgt $\mu_1^{(i)} = \dots = \mu_{r_i}^{(i)} = 0$.

ii) $\Leftrightarrow$ iii) ist klar.

ii) $\Rightarrow$ i). Zum Nachweis von DS2 stellen wir jedes $w_i \in W_i$ durch die gegebene Basis von W_i dar:

$$w_i = \mu_1^{(i)} v_1^{(i)} + \dots + \mu_{r_i}^{(i)} v_{r_i}^{(i)}.$$

Aus der Annahme $w_1 + \dots + w_k = 0$ folgt

$$\sum_{i=1}^k \sum_{\varrho=1}^{r_i} \mu_{\varrho}^{(i)} v_{\varrho}^{(i)} = 0.$$

Da $\mathcal{B}$ eine Basis ist, folgt $\mu_{\varrho}^{(i)} = 0$ für alle ϱ und i , also ist auch

$$w_1 = \dots = w_k = 0.$$

□

Aufgaben zu 1.6

1. Beweisen Sie, dass für einen Vektorraum V folgende Bedingungen äquivalent sind:

i) $V = W_1 \oplus \dots \oplus W_k$.

ii) Jedes $v \in V$ ist eindeutig darstellbar als $v = w_1 + \dots + w_k$ mit $w_i \in W_i$.

iii) $V = W_1 + \dots + W_k$, $W_i \neq \{0\}$ für alle i und von Null verschiedene Vektoren $w_1 \in W_1, \dots, w_k \in W_k$ sind linear unabhängig.

Vorsicht! Die Voraussetzung $W_i \neq \{0\}$ ist wesentlich. Ist etwa $W_1 = \{0\}$, so ist die angegebene zweite Bedingung stets erfüllt!

iv) $V = W_1 + \dots + W_k$ und $W_i \cap \sum_{\substack{j=1 \\ j \neq i}}^k W_j = \{0\}$ für alle $i \in \{1, \dots, k\}$.

v) $V = W_1 + \dots + W_k$ und $W_i \cap (W_{i+1} + \dots + W_k) = \{0\}$ für alle $i \in \{1, \dots, k-1\}$.

Zeigen Sie anhand von Gegenbeispielen, dass die obigen Bedingungen für $k > 2$ im Allgemeinen nicht äquivalent sind zu $W_1 \cap \dots \cap W_k = \{0\}$ bzw. $W_i \cap W_j = \{0\}$ für alle $i \neq j$.

2. Sind V und W Vektorräume, so gilt

$$V \times W = (V \times \{0\}) \oplus (\{0\} \times W).$$

3. Eine Matrix $A \in M(n \times n; K)$ heißt *symmetrisch*, falls $A = {}^t A$.

a) Zeigen Sie, dass die symmetrischen Matrizen einen Untervektorraum $\text{Sym}(n; K)$ von $M(n \times n; K)$ bilden. Geben Sie die Dimension und eine Basis von $\text{Sym}(n; K)$ an.

Ist $\text{char } K \neq 2$, so heißt $A \in M(n \times n; K)$ *schiefssymmetrisch* (oder *alternierend*), falls ${}^tA = -A$. Im Folgenden sei stets $\text{char } K \neq 2$.

- b) Zeigen Sie, dass die alternierenden Matrizen einen Untervektorraum $\text{Alt}(n; K)$ von $M(n \times n; K)$ bilden. Bestimmen Sie auch für $\text{Alt}(n; K)$ die Dimension und eine Basis.
- c) Für $A \in M(n \times n; K)$ sei $A_s := \frac{1}{2}(A + {}^tA)$ und $A_a := \frac{1}{2}(A - {}^tA)$. Zeigen Sie: A_s ist symmetrisch, A_a ist alternierend, und es gilt $A = A_s + A_a$.
- d) Es gilt: $M(n \times n; K) = \text{Sym}(n; K) \oplus \text{Alt}(n; K)$.