



UNIVERSITÄT
HEIDELBERG
ZUKUNFT
SEIT 1386

SKRIPT ZUR VORLESUNG
HÖHERE MATHEMATIK I
(FÜR PHYSIKSTUDIERENDE)

gehalten an der
Universität Heidelberg
im
Wintersemester 2024/25
von
Johannes Walcher

INHALTSVERZEICHNIS

1	GRUNDLAGEN	3
§ 1	Logik	4
§ 2	Mengen	8
§ 3	Relationen	11
§ 4	Abbildungen	12
2	ZAHLEN	19
§ 5	Gruppen, Ringe, Körper	20
§ 6	Die reellen Zahlen	24
§ 7	Die komplexen Zahlen	32
§ 8	Physikalische Räume	35
3	VEKTOREN	43
§ 9	Vektorräume	44
§ 10	Basen, Matrizen, Determinanten	56
§ 11	Lineare Abbildungen	72
§ 12	Summen-, Quotienten-, Dualräume	83
§ 13	Lineare Gleichungen	91
4	KONVERGENZ	95
§ 14	Folgen reeller Zahlen	96
§ 15	Metrische Räume	110
§ 16	Reihen	123

KAPITEL 1

GRUNDLAGEN

Die Vorlesungsreihe “Höhere Mathematik” richtet sich an Studierende der Physik in den ersten drei Semestern. Ihre Aufgabe ist die Vermittlung der für das Studium der Physik essentiell notwendigen mathematischen Konzepte und Rechenmethoden. Sie umfasst (in ihren Grenzen) in etwa den gleichen Stoff wie die fünf(!) Grundmodule “Lineare Algebra” und “Analysis” aus dem Bachelorstudiengang Mathematik. Die zeitliche Straffung wird nicht durch prinzipiellen Verzicht auf mathematischen Tiefgang und Strenge erreicht, sondern durch das Ersetzen einiger Details durch Verweise auf das reichhaltige Anschauungsmaterial aus den physikalischen Kursvorlesungen. Dies betrifft als Beispiel insbesondere explizite Lösungsmethoden für gewöhnliche Differentialgleichungen oder die geometrisch-physikalische Interpretation von Integralsätzen. In Folge dessen fehlen nicht nur Beweistechniken sondern auch die intrinsische mathematische Motivation und Begriffsbildung. Daher ist die HöMa nur bedingt als Vorbereitung auf weiterführende mathematische Vorlesungen geeignet. Studierenden, die einen stark formal-theoretischen Einschlag verspüren oder womöglich noch zwischen Physik und Mathematik als Studienfach schwanken, wird empfohlen, den Besuch der LA/Ana in Betracht zu ziehen.

Das erste Semester beginnt mit einigen Grundgedanken zur mathematischen Sprache und Notation sowie zur elementaren Mengenlehre. Im Anschluss an die Vorstellung der wichtigsten Zahlbereiche als Grundlage für die quantitative Erfassung von Naturphänomenen untersuchen wir zunächst die lineare Struktur von (endlich-dimensionalen) Vektorräumen, linearen Abbildungen und die Lösung linearer Gleichungssysteme. In der zweiten Semesterhälfte geht es dann vor allem um die Entwicklung des Konvergenz- und Vollständigkeitsbegriffes, insbesondere im Setting normierter Vektorräume. Wir beweisen damit den Fundamentalsatz der Algebra, bevor wir im letzten Abschnitt mit einer Einführung in die Eigenwerttheorie zur linearen Algebra zurückkehren.¹

· Für Kommentare, auch Fehlermeldungen, zum Skript, per Email an walcher@uni-heidelberg.de, bin ich sehr dankbar.

· Homepage der Vorlesung im Wintersemester 2024/25:
web.mathi.uni-heidelberg.de/physmath/wise2425/hoema1

¹Ein kurzer Blick in das Inhaltsverzeichnis offenbart, dass wir dieses letzte Ziel nicht erreicht haben werden.

§ 1 Logik

In der ursprünglichen Bedeutung des Wortes bezeichnet *Mathematik* (was ich gerne übersetze als) das “Handwerk zur Einsicht”. Mit dieser (meiner) Übersetzung soll betont werden, dass “Verstehen” eine abstrakt-geistige Leistung ist, und sich durch konkrete Praxis erwerben lernen lässt.

Etwas präziser und moderner ist das Geschäft der Mathematik, aus gewissen als “wahr” angenommenen Aussagen (den “Axiomen”) die Wahrheit anderer Aussagen (den “Theoremen”) über präzis definierte “mathematische Objekte” logisch herzuleiten (zu “beweisen”). Der Wert der Mathematik (*sc.* für die Physik) ergibt sich aus dem *Ummünzen* derartiger Aussagen zu konkreten Aussagen über natürliche Sachverhalte, d.h. zur Naturbeschreibung.

Was diese Aussagen genau bedeuten, lässt sich wie eingangs angedeutet nur langsam und allmählich durch Betrachten von Beispielen erfassen. Dies liegt in ganz charakteristischer Weise daran, dass wir ausser durch Vorzeigen gar nicht genau erklären können, was mit den einzelnen Vokabeln und dem Satz als Ganzes (und noch weniger mit dem Bezug zur Natur) eigentlich gemeint ist. Es ist aber hilfreich, dass wir uns, bevor wir uns an das eigentliche Geschäft machen, wenigstens ansatzweise über die Umrisse dieser Bedeutungen *einigen*. Dies ist der Zweck der folgenden “Vereinbarungen”.

Vereinbarung 1.1. Eine Aussage ist ein feststellender Satz, dem in sinnvoller und eindeutiger Weise der Wert “wahr” (*w*) oder “falsch” (*f*) zugeordnet ist, sein oder werden kann.

Beispiele. “Es regnet.”; “Donald J. Trump is President of the United States.” sind Aussagen. “Sick!” ist keine Aussage.

Eine Aussage ist demnach zunächst einmal eine nach den syntaktischen Regeln einer Sprache gebildete Folge von Zeichen. Diese Zeichen können visuell, akustisch, oder sonstwie irgendeiner Intelligenz verständlich sein. Sie können elementar oder zusammengesetzt sein. Welche Zeichen erlaubt sind und welche Regeln zu beachten sind (in der Mathematik gelten ganz besondere), ergibt sich aus dem sprachlichen Kontext. Allerdings ist auch nicht jeder den Regeln der Grammatik genügender Satz unabhängig vom Kontext notwendiger Weise eine Aussage. Darüber hinaus wird in dieser Vereinbarung nicht spezifiziert, von wem, wann und wie über den Wahrheitswert einer Aussage entschieden wird.

Beispiele. “Die Zahl 7 ist verheiratet.”; “Gott ist tot.”; “Es wird morgen um 10h00 MESZ im Heidelberger Stadtgebiet regnen.”

Für das Geschäft der mathematischen Logik (das man formal ein Leben lang betreiben kann) ist es wichtig, dass wir

- * über Aussagen verständig reden können (was wir gerade tun);
- * Aussagen symbolisieren können (z.B. A: “Der Prof. schreibt auf einem iPad”; B: “Apple ist die Tochter von Gwyneth Paltrow.”);
- * gegebene Aussagen zu neuen Aussagen verknüpfen können.

Wir erklären solche Verknüpfungen mit Hilfe sogenannter *Wahrheitstafeln*, und zwar:

§1. LOGIK

· die Negation (logisches “nicht”):	A	$\neg A$
	w	f
	f	w

· die Konjunktion (logisches “und”):	A	B	$A \wedge B$
	w	w	w
	w	f	f
	f	w	f
	f	f	f

· die Disjunktion (logisches “oder”):	A	B	$A \vee B$
	w	w	w
	w	f	w
	f	w	w
	f	f	f

Übungsaufgabe. Die Aussage $A \vee B$ ist logisch gleichwertig, d.h. hat den gleichen Wahrheitswert wie, die Aussage $\neg(\neg A \wedge \neg B)$ — Dass A oder B wahr ist heisst so viel, wie dass A und B nicht beide falsch sind. Ebenso ist $\neg\neg A$ gleichwertig zu A (“tertium non datur”).

Besonders wichtig und gewöhnungsbedürftiger als der anfängliche Schein ist

· die Implikation (logisches “wenn-dann”):	A	B	$A \Rightarrow B$
	w	w	w
	w	f	f
	f	w	w
	f	f	w

Hierbei heisst A das Antezedens und B das Konsequens (der Implikation $A \Rightarrow B$). Im mathematischen Kontext interpretiert man vor allem A als *hinreichende Bedingung* für B , und umgekehrt B als *notwendige Bedingung* für A . Beachte allerdings, dass der Wahrheitswert einer Wenn-Dann-Implikation nichts über eine materielle Ursache-Wirkung-Beziehung zwischen A und B aussagt (s. Beispiel).

Zur Belegung der zwei letzten Zeilen (“ex falso quodlibet”): Dass $A \Rightarrow B$ wahr ist, sobald A falsch ist, ist letztlich eine Konvention, die sich in der Praxis bewährt hat. Sie kann z.B. dadurch als sinnvoll eingesehen werden, dass man die Aussage “Wenn A dann B ” nicht dadurch widerlegen können möchte, dass man ein falsches A zusammen mit einem beliebigen (wahren oder falschen) B vorzeigt, andererseits Aussagen aber einen eindeutigen Wahrheitswert brauchen.

Ähnliche Bemerkungen gelten für

· die Äquivalenz (logisches “dann-und-nur-dann” bzw. “genau-dann-wenn”):	A	B	$A \Leftrightarrow B$
	w	w	w
	w	f	f
	f	w	f
	f	f	w

Die obigen Verknüpfungen bewerten Beziehungen zwischen prinzipiell fest stehenden oder gedanklich fest gehaltenen Aussagen, deren Wahrheitswert bestenfalls kontextuell ist. Daher lässt sich daraus schwerlich neue Erkenntnis gewinnen. Interessanter und ergiebiger wird es wenn wir Aussagen über quantifizierte Gesamtheiten

von Objekten zulassen. Hierzu benötigen wir logische Variablen, Prädikate (Aussagen mit Lücken) und die sogenannten Quantoren. Wir symbolisieren:

- * Variablen als Platzhalter oder Zeichen für Objekte mit Buchstaben x, y , etc.;
- * Prädikate mit einer Leerstelle oder Argument mit $A(\cdot)$, zwei Leerstellen $B(\cdot, \cdot)$, usw. Sie werden zu Aussagen (“Formeln erster Ordnung”) durch Einsetzen von Variablen $A(x)$, $B(x, y)$, usw.
- * Den Universalquantor als umgedrehtes ‘A’: $\forall$; und den Existenzquantor als umgedrehtes ‘E’: $\exists$

und erklären:

- die Aussage $\forall x:A(x)$ ist wahr wenn für alle x die Aussage $A(x)$ wahr ist. Sonst ist $\forall x:A(x)$ falsch.
- die Aussage $\exists x:A(x)$ ist wahr, wenn ein x existiert, für welches die Aussage $A(x)$ wahr ist. Sonst ist $\exists x:A(x)$ falsch.

Weiterhin bedeute $\exists!x:A(x)$, dass genau ein x existiert so, dass $A(x)$ wahr ist. Dies ist äquivalent zu $\exists x:(A(x) \wedge (\forall y:A(y) \Rightarrow x = y))$. Hierbei symbolisiert $x = y$ natürlich die “Gleichheit” der bezeichneten Objekte.

Solange wir Aussagen nur als “formale Derivate” (be)handeln, denen wir vielleicht prinzipiell Wahrheitswerte zuordnen können (aber nicht unbedingt müssen oder wollen), ist es relativ unerheblich, was mit Existenz bzw. Universalität genau gemeint ist. In der Praxis, d.h. um über Wahrheit/Falschheit tatsächlich zu entscheiden, kommt es natürlich essentiell darauf an. Man nennt dabei den Bereich, über den x, y , usw. variieren das *Diskursuniversum* (oder die Grundmenge). Sehr bald kann man dann allerdings über die Wahrheit von Aussagen nicht mehr mit Wahrheitstabellen entscheiden.

Beispiele. “Es hat jemand in diesem Raum am 29. Februar Geburtstag.”; “Alle Menschen (*sc.* die aktuell auf der Erde leben/bisher gelebt haben) sind sterblich.”; “Es existiert eine grösste Primzahl.”; “Es gibt nur einen Einstein.”

Wir unterbrechen nun an dieser Stelle diese philosophisch-psychologische Diskussion und machen stattdessen noch einige Bemerkungen dazu, wie wir kompliziertere Aussagen und Objekte konstruieren und damit umgehen wollen.

Vereinbarung 1.2. Unter einer Definition verstehen wir (primär/meistens) die Reservierung von Symbolen und Vokablen als Abkürzung von Aussagen (Formeln erster Ordnung) zur Erweiterung der (mathematischen) Sprache.

Wir formulieren solche Definitionen mit Hilfe von Sätzen der natürlichen Sprache, der Form “Ein Objekt x heisst __blah__ falls die Aussage $A(x)$ wahr ist.” (bzw. einfach nur: “falls $A(x)$ ”) Hiermit ist implizit mitgemeint, dass falls $A(x)$ falsch ist, x nicht __blah__ heisst. Wichtig ist zur “Wohldefiniiertheit”, dass die verwendeten Aussagen tatsächlich zu wahr oder falsch ausgewertet werden können. (Vgl. die allgemeinen Bemerkungen zu Aussagen oben.)

Beispiel. Eine Primzahl ist eine natürliche Zahl, die durch genau eine andere Zahl (nämlich die 1) ohne Rest teilbar ist. (Mit “nur sich selbst und die 1” statt “genau eine andere Zahl“ wäre die 1 selbst auch eine Primzahl, was man aber nicht will.)

§ 1. LOGIK

Wichtig wäre ausserdem, dass (wenigstens in den Grundbegriffen und innerhalb von mathematischen Spezialgebieten) die Bedeutung von mathematischen Vokabeln ein-eindeutig ist, d.h. zu jedem Konzept gibt es nur genau eine Vokabel und jede Vokabel bezeichnet nur genau ein Konzept. Verletzungen dieser Regel, die historisch und soziologisch unvermeidbar sind, im Zuge besserer Kommunikation aber etwas harmonisiert worden sind, signalisiert man durch den Ausruf “abuse of language!”.

Unfreiwilliges Beispiel hierfür ist der Begriff der Definition selber, auch wenn er wohlgemerkt “nur” eine Vereinbarung ist und keine Definition an sich. Denn häufig/sekundär bezeichnet der Begriff der Definition auch die “Konstruktion” bzw. das “Hervorzeigen” konkreter mathematischer Objekte. Hier ist zur Wohldefiniertheit notwendig, dass die (im obigen primären Sinne) definierenden Eigenschaften auch tatsächlich erfüllt sind. Als Symbol benutzen wir in solchen Definitionen manchmal das Kolon ‘:’, aber auch dies bedeutet häufig etwas anderes (z.B., “so, dass” oder “es gilt”).

Beispiel. Das aus der Schule bekannte “Gegeben ist die Funktion f mit $f(x) = x^2$ ” werden wir nicht als (vollständige) Definition in unserem Sinne zulassen.

Der eigentliche Akt des mathematischen Geschäfts ist es dann, mit Hilfe solcher Definitionen die “tautologische” Wahrheit von (komplizierten) Aussagen zu behaupten und anschliessend zu beweisen. Dabei nennen wir solche Aussagen, in Abhängigkeit der Wichtigkeit für die weitere Entwicklung des Stoffs und der Schwierigkeit des Beweises “Lemma”, “Proposition”, “Theorem” oder “Korollar”. Typische Form solcher Aussagen sind Implikationen $A \Rightarrow B$, deren Wahrheit mit Hilfe von Wahrheitstabellen nicht oder nur schwer einzusehen ist. Hierbei heisst A Voraussetzung und B Folgerung der Aussage des Lemmas oder Theorems.

Vereinbarung 1.3. Ein Beweis ist eine Kette von Aussagen (der mathematischen Sprache) deren Wahrheitswerte mit Hilfe einer möglicherweise erweiterten natürlichen Sprache miteinander verknüpft sind und aus der die notwendige Wahrheit von B aus der von A hervorgeht.

Auch hier werden/müssen wir anstatt endlos zu philosophieren oder ausgiebig über Beweistechniken zu theoretisieren hauptsächlich die Beweise selber sprechen lassen. Eine wichtige Einsicht, die wir allerdings bereits hier festhalten wollen, ist dass aufgrund der logischen Äquivalenz von $A \Rightarrow B$ mit $\neg B \Rightarrow \neg A$ man eine Implikation auch beweisen kann, indem man die Folgerung negiert und daraus die Negation der Voraussetzung herleitet. Diese sog. Kontraposition ist der wichtigste Spezialfall des Widerspruchsbeweises (“reductio ad absurdum”).

Für das eigentliche logische Schliessen werden wir keine eigene Notation verwenden. (Typischerweise zählen hierzu: $\therefore$ für “deshalb”, $\models$ für “systematische/formale Implikation”, und $\vdash$ für “Beweis-theoretische Implikation”.) Einzige Ausnahmen:

- “quod erat demonstrandum” (was zu beweisen war), für “Der Beweis ist erbracht, ab jetzt können wir die Aussage (die behauptete tautologische Implikation) als wahr betrachten und weiterverwenden.”
- ⚡ als Zeichen für “dies widerspricht der Voraussetzung”, woraus dann (in einem Beweis durch Kontraposition) üblicherweise □ folgt.

Zuletzt sei noch bemerkt, dass sich Beweise technisch automatisieren (und dann von ChatGPT in die natürliche Sprache zurückübersetzen) lassen. Ob sich damit allerdings echte Einsicht oder neue Erkenntnisse gewinnen lassen, sei dahin gestellt. Jedenfalls in dieser Vorlesung wird über die Wahrheit von Aussagen und über die Richtigkeit von Beweisen im gemeinsamen Gespräch (bzw. in Ihrer Auseinandersetzung mit den Tutoren) entschieden.

Über die Grenzen der Logik reden wir nicht.

§ 2 Mengen

Die Mengenlehre ist eine in den Ansätzen recht einfache mathematische Theorie, die sich—trotz einiger wohlbekannten Unzulänglichkeiten—als Grundlage für den gesamten weiteren Aufbau bewährt hat.² Vieles davon ist intuitiv leicht zu erfassen oder bereits aus der Schule bekannt. Deshalb geht es auch vor allem um die Festlegung von Begriffen und der Notation.³

Vereinbarung 2.1. (Cantor, 1895)

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

Dies ist der letzte Begriff, dessen Bedeutung wir nicht durch eine Definition festlegen. Als (vorläufig) letzte philosophische Bemerkung sei auf die Wichtigkeit des Einschubs “unsrer Anschauung oder unseres Denkens” für die Physik hingewiesen: Erst die Gleichstellung des Konkreten mit dem Abstrakten innerhalb eines gemeinsamen Rahmens erlaubt nämlich den präzisen Vergleich der beiden Sphären und damit überhaupt die quantitative Naturbeschreibung. So beruht z.B. die Aussage “Unser Sonnensystem hat 8 Planeten.” auf einer Korrespondenz zwischen der anschaulichen Menge der Planeten und der gedachten Menge der ersten 8 natürlichen Zahlen.

Wir benutzen nun (für eine kleine Weile) Grossbuchstaben A, B , etc. für Mengen und Kleinbuchstaben $a, b, c, \dots$ für (potentielle) Elemente und legen fest, dass

- * wir schreiben $a \in A$ für die Aussage “ a ist Element von A ”, und $a \notin A$ für $\neg(a \in A)$;
- * wir können Mengen definieren (im Sinne von “angeben”), und zwar
 - entweder *extensiv* durch Angabe ihrer Elemente, z.B. $A := \{1, 2, 3, 4, 5, 6, 7, 8\}$,
 $B := \{\text{Merkur, Venus, Erde, Mars, Jupiter, Saturn, Uranus, Neptun}\}$
 - oder *intensiv* durch eine charakterisierende Eigenschaft, z.B. $C := \{a \mid E(a)\}$,
 wo E ein Prädikat ist, das innerhalb eines gegebenen Diskursuniversums Sinn macht.

²Es gibt allerdings auch andere Möglichkeiten zur Fundierung der Mathematik als die Mengenlehre. Einige Stichworte: Synthetische Geometrie, Kategorientheorie, Typentheorie.

³Wie auch bei den später behandelten Theorien ist unsere Vorstellung der Mengenlehre natürlich nicht sofort vollständig und wird bei Bedarf ergänzt.

§ 2. MENGEN

* eine Menge $\emptyset$ existiert (die leere Menge), mit der Eigenschaft, dass $\forall a : a \notin \emptyset$.

Man schreibt auch $\emptyset = \{\}$.

Sodann können wir loslegen.

Definition 2.2. (i) Eine Menge A heisst *Teilmenge* einer Menge B (geschrieben: $A \subset B$) falls $\forall a : a \in A \Rightarrow a \in B$.

(ii) A und B heissen gleich (geschrieben $A = B$) falls $A \subset B \wedge B \subset A$.

(iii) A heisst *echte Teilmenge* von B (geschrieben: $A \subsetneq B$, $A \subsetneqq B$, $A \subsetneq B$, oder $A \subsetneq B$) falls $A \subset B$ aber $B \not\subset A$.⁴

Beispiel. $\{1, 2\} = \{2, 1\} = \{1, 1, 2\}$, d.h. es kommt bei der extensiven Definition nicht auf die Reihenfolge der Elemente an, und Mehrfachnennung von Elementen ändert die Menge nicht.

Lemma 2.3. (i) Es seien A, B, C Mengen mit der Eigenschaft, dass $A \subset B$ und $B \subset C$. Dann gilt $A \subset C$.

(ii) Für jede Menge A gilt: $\emptyset \subset A$.

(iii) Für jede Menge N mit der Eigenschaft $\forall a : a \notin N$ gilt: $N = \emptyset$.

In Worten: (i) Mengeninklusion ist transitiv (s. Definition 3.2). (ii) Jede Menge besitzt (ausser sich selbst) immer die leere Menge als Teilmenge. (Dies gilt auch für die leere Menge selber!) (iii) Es gibt nur eine leere Menge.

Beweis von Lemma 2.3. (i) Wir müssen zeigen: $\forall a : a \in A \Rightarrow a \in C$. Nach Voraussetzung $A \subset B$ folgt aber aus $a \in A$, dass auch $a \in B$, und nach Voraussetzung $B \subset C$ folgt daraus $a \in C$. Daher ist $a \in A \Rightarrow a \in C$ unter den gegebenen Voraussetzungen stets wahr.

(ii) Die zu zeigende Aussage: $\forall a : a \in \emptyset \Rightarrow a \in A$, die zunächst etwas seltsam anmutet folgt direkt aus der vormals etwas seltsam anmutenden Wahrheitstafel für die Implikation (S. 5): $a \in \emptyset$ ist immer falsch, und daher die Implikation stets wahr.

(iii) Funktioniert genauso. $\square$

Definition 2.4. Sind A und B Mengen, so heissen

(i) $A \cap B := \{a \mid a \in A \wedge a \in B\}$ der *Durchschnitt* von A und B ; A und B heissen *disjunkt*, falls $A \cap B = \emptyset$;

(ii) $A \cup B := \{a \mid a \in A \vee a \in B\}$ die *Vereinigung* von A und B ; gilt ausserdem $A \cap B = \emptyset$, so nennt man $A \cup B$ *disjunkte Vereinigung* und man schreibt zur Betonung $A \dot{\cup} B$ oder $A \sqcup B$;

(iii) $A \setminus B := \{a \mid a \in A \wedge a \notin B\}$ das *Komplement* von B in A ;

(iv) $A \times B := \{(a, b) \mid a \in A, b \in B\}$ das *kartesische Produkt* von A und B , wobei (a, b) das geordnete Paar von a und b bezeichnet.

Wir veranschaulichen diese Operationen durch sog. Venn-Diagramme und benutzen ohne viel Aufhebens Rechenregeln der Art:

$$* A \cap B \subset B \subset A \cup B$$

⁴In manchen Zusammenhängen ist es sinnvoll, die leere Menge (die gemäss Lemma 2.3 Teilmenge einer jeden Menge ist,) nicht als "echte" Teilmenge zu bezeichnen.

$$* (A \cap B) \cup C = (A \cup C) \cap (B \cup C)$$

$$* C \setminus (A \cup B) = (C \setminus A) \cap (C \setminus B)$$

$$\vdots$$

(die Liste ist stark unvollständig!)

$$\vdots$$

die man sich bei Bedarf leicht “vor Ort” klar macht. Zum kartesischen Produkt sind allerdings zwei Dinge zu bemerken.

1. Es kommt auf die Reihenfolge an, d.h. $A \times B \neq B \times A$ (das kartesische Produkt ist nicht kommutativ)

2. Es gibt eine Subtilität bei seiner mehrfachen Anwendung. Meistens (aber nicht immer!) wird vereinbart, dass $(A \times B) \times C = \{(a, b), c\} = \{(a, b, c)\}$ bedeutet, d.h. die runden Klammern bedeuten “Wortbildung” bzw. “Wortfortsetzung”. Dann ist das kartesische Produkt assoziativ, d.h. $(A \times B) \times C = A \times (B \times C)$.

Unsere Vereinbarung 2.1 enthält keine Einschränkungen an die Sorte Objekte die in einer Menge enthalten sein können. Insbesondere können Mengen auch selbst wieder Elemente von Mengen sein. (Dabei ist $\{\{1, 2\}, 3\} \neq \{1, \{2, 3\}\}$.) In letzter Konsequenz führt dies zu Konflikten mit den übrigen als sinnvoll erachteten Axiomen (Stichwort: Russelsche Antinomie), was aber für die Praxis unerheblich ist.

Definition 2.5. Sei A eine Menge. Dann heisst die Menge aller Teilmengen von A die *Potenzmenge von A* , geschrieben

$$\mathcal{P}(A) := \{B \mid B \subset A\} \quad (2.1)$$

Lemma 2.6. Es gilt $\emptyset \in \mathcal{P}(A)$ und $A \in \mathcal{P}(A)$ für jede Menge A . $\mathcal{P}(\emptyset) = \{\emptyset\} \neq \emptyset$.

Die Menge der natürlichen Zahlen wird als bekannt vorausgesetzt und mit $\mathbb{N}$ bezeichnet. Man schreibt üblicherweise $\mathbb{N} = \{1, 2, 3, \dots\}$, was aber streng genommen keine Definition ist. Charakteristisch ist vielmehr das *Prinzip der vollständigen Induktion*:

Eine universelle Aussage $\forall n : n \in \mathbb{N} \wedge A(n)$ über die Menge der natürlichen Zahlen (man schreibt auch $\forall n \in \mathbb{N} : A(n)$) ist genau dann wahr, wenn

der *Induktionsanfang*: $A(1)$ wahr ist

und

der *Induktionsschritt*: $\forall n \in \mathbb{N} : A(n) \Rightarrow A(n+1)$ wahr ist.

Anders ausgedrückt ist eine natürliche Zahl vollständig und eindeutig definiert über die Menge ihrer Vorläufer. Deshalb durchläuft man alle ohne Wiederkehr, indem man bei der 1 startet und bei jeder (nach “Geniessen der Aussicht”) von dieser zu ihrer Nachfolgerin weiterschreitet.

Beispiel 2.7. Für alle $n \in \mathbb{N}$ gilt $\sum_{k=1}^n k := 1 + 2 + \dots + n = \frac{n(n+1)}{2}$.

Induktionsanfang: Für $n = 1$ gilt $\sum_{k=1}^1 k = 1 = \frac{1 \cdot 2}{2}$.

Induktionsschritt: Aus $1 + \dots + n = \frac{n(n+1)}{2}$ folgt

$$1 + \dots + n + (n+1) = \frac{n(n+1)}{2} + n+1 = \frac{n(n+1) + 2(n+1)}{2} = \frac{(n+1)(n+2)}{2} \quad (2.2)$$

□

§ 3. RELATIONEN

Viele andere Mengen der Mathematik lassen sich aus der Menge der natürlichen Zahlen explizit konstruieren.

$$\cdot \mathbb{N}_0 := \mathbb{N} \cup \{0\}$$

$$\cdot \mathbb{Z} := \{0\} \cup (\{+, -\} \times \mathbb{N}) = \{0, 1, -1, 2, -2, \dots\}$$

$\vdots$

Wenn das Alphabet nicht ausreicht, benutzen wir zur Unterscheidung von Variablen, Mengen etc. gerne die natürlichen Zahl als tief- und manchmal auch hochgestellte *Indizes*. Bsp.: $A_1, A_2, \dots; x^1, x^2, \dots$ (Die Möglichkeit der Verwechslung von Superskripten mit Potenzen ist selten gefährlich.) Dies wird in den folgenden Kapiteln noch weiter gedeutet und formalisiert.

§ 3 Relationen

Beziehungen zwischen verschiedenen (mathematischen und physikalischen) Objekten lassen sich im Rahmen der Mengenlehre wie folgt formalisieren.

Definition 3.1. Eine (zweistellige oder binäre) *Relation* ist eine Teilmenge $R \subset A \times B$ des kartesischen Produkts zweier Mengen A und B .

Wir notieren Relationen typischerweise als: aRb falls $(a, b) \in R$ und veranschaulichen sie mittels gerichteter Graphen. Wir sprechen zunächst über Relationen auf nur einer Menge, d.h. der Fall $A = B$.

Definition 3.2. Eine Relation $R \subset A \times A$ heisst

- (i) *reflexiv*, falls aRa für alle $a \in A$;⁵
- (ii) *irreflexiv*, falls aRa für kein $a \in A$;
- (iii) *symmetrisch*, falls $a_1Ra_2 \Rightarrow a_2Ra_1$ für alle $a_1 \in A$ und $a_2 \in A$;
- (iv) *anti-symmetrisch*, falls $a_1Ra_2 \wedge a_2Ra_1 \Rightarrow a_1 = a_2$;⁶
- (v) *transitiv*, falls $a_1Ra_2 \wedge a_2Ra_3 \Rightarrow a_1Ra_3$;
- (vi) *total*, falls für alle $a_1, a_2 \in A$ a_1Ra_2 oder a_2Ra_1 wahr ist.
- (vii) *trichotomisch*, falls für alle $a_1, a_2 \in A$ genau eine der drei Aussagen a_1Ra_2 , a_2Ra_1 oder $a_1 = a_2$ wahr ist.

Diese Vokabeln bedingen sich gegenseitig in nicht ganz offensichtlicher Weise. So ist z.B. eine totale Relation automatisch reflexiv und eine trichotomische Relation automatisch irreflexiv. Für uns spielen allerdings nur zwei Spezialfälle eine wichtige Rolle.

Definition 3.3. Eine Relation R auf A heisst

- (i) *partielle Ordnung*, falls R reflexiv, anti-symmetrisch und transitiv ist.
- (ii) *Totalordnung*, falls R anti-symmetrisch, transitiv und total ist.
- (iii) *strenge Totalordnung*, falls R transitiv und trichotomisch ist.

⁵Wir richten uns hier wie meistens doch an dem Usus aus, Quantoren hintanzustellen und auszuschreiben...

⁶...und manchmal auch ganz wegzulassen

In solchen Fällen nennt man das Paar (A, R) ⁷ eine (partiell, total oder streng total) (an-)geordnete Menge. Wir schreiben im Allgemeinen partielle und Totalordnungen als $\preceq$ und strenge Totalordnungen als $\prec$

Übungsaufgabe. Jede Totalordnung ist auch eine partielle Ordnung. Eine strenge Totalordnung ist zwar keine Totalordnung, induziert aber eine via $a \preceq b :\Leftrightarrow (a \prec b \vee a = b)$

Beispiele 3.4. * Auf der Menge der natürlichen Zahlen ist die “kleiner-als”-Relation $R = \{(n_1, n_2) \in \mathbb{N} \times \mathbb{N} \mid n_1 < n_2\}$ eine strenge Totalordnung.

* Auf der Menge der natürlichen Zahlen ist die “kleiner-gleich”-Relation $R = \{(n_1, n_2) \mid n_1 < n_2 \text{ oder } n_1 = n_2\}$ eine Totalordnung.

* Auf der Potenzmenge $\mathcal{P}(A)$ einer jeden Menge A ist die Mengeninklusion eine partielle Ordnung (Lemma 2.3 zeigt die Transitivität.)

Ordnungen sind also intuitiv nichts anderes als eine Formalisierung von Ungleichungen. Wie steht es mit Gleichungen?

Definition 3.5. Eine Relation R auf einer Menge A heisst *Äquivalenzrelation*, falls R reflexiv, symmetrisch und transitiv ist. Wir schreiben im Allgemeinen Äquivalenzrelationen als $\sim$. Eine Teilmenge $C \subset A$ heisst *Äquivalenzklasse bezüglich $\sim$* , falls (i) $C \neq \emptyset$; (ii) für alle $a, b \in C$ gilt, dass $a \sim b$ und (iii) aus $a \in C$ und $b \sim a$ folgt, dass $b \in C$. Die Elemente einer Äquivalenzklasse heissen ihre *Repräsentanten*.

Beispiele 3.6. * Auf (der Menge der locations auf) dem Königstuhl ist die Relation $loc_1 \sim loc_2$ falls loc_1 und loc_2 auf der gleichen Höhe über dem Meer liegen, eine Äquivalenzrelation. Äquivalenzklassen sind die Höhenlinien.

* Sei $n \in \mathbb{N}$. Dann definiert die Bedingung

$$k_1 \sim k_2 \Leftrightarrow n \mid k_1 - k_2 \text{ (dies heisst: } k_1 - k_2 \text{ ist durch } n \text{ ohne Rest teilbar)} \quad (3.1)$$

eine Äquivalenzrelation auf der Menge $\mathbb{Z} \ni k_1, k_2$ der ganzen Zahlen, genannt *Kongruenz modulo n* , geschrieben $k_1 \equiv_n k_2$. Die Äquivalenzklassen heissen *Restklassen modulo n* .

* Auf der Menge von Paaren $(p, q) \in \mathbb{Z} \times (\mathbb{Z} \setminus 0)$ wird durch

$$(p_1, q_1) \sim (p_2, q_2) \Leftrightarrow p_1 q_2 = p_2 q_1 \quad (3.2)$$

eine Äquivalenzrelation definiert. Äquivalenzklassen sind *rationale Zahlen* (s.u.)

§ 4 Abbildungen

Unter den Relationen zwischen zwei *verschiedenen* Mengen sind solche von besonderem Interesse, die es erlauben, Objekte verschiedener Art durcheinander zu *ersetzen*. Aus notationstaktischen Gründen bezeichnen wir Mengen nun mit Grossbuchstaben vom Ende des Alphabets.

⁷Ein solches Paar sollte ein Element der Menge $\{(A, \mathcal{P}(A \times A)) \mid A \text{ ist eine Menge}\}$ sein. Diese wiederum wäre sowohl ein Element der “Menge aller Mengen” $\mathcal{M}$, als auch eine Teilmenge von $\mathcal{M} \times \mathcal{M}$, und es gälte $\mathcal{M} \times \mathcal{M} \subsetneq \mathcal{M}$ — Dies zur Illustration der Idee “Alles ist eine Menge” und ihrer abenteuerlichen Konsequenzen.

§4. ABBILDUNGEN

Definition 4.1. Ein Tripel $f = (X, Y, G)$ von Mengen mit der Eigenschaft, dass

- (i) $G \subset X \times Y$ (d.h. G ist eine Relation zwischen X und Y) und
- (ii) für alle $x \in X$ existiert genau ein $y \in Y$ so, dass $(x, y) \in G$

heißt *Abbildung von X nach Y* . Dabei heißt X der *Definitionsbereich*, Y der *Wertebereich* und G der *Graph* der Abbildung. Man schreibt für das durch x eindeutig bestimmte Element von Y normalerweise $f(x)$ und sagt $x \in X$ werde auf $f(x) \in Y$ *abgebildet*. Hierfür schreibt man auch $x \mapsto f(x)$ und nennt $f(x)$ das “Bild von x unter f ”. Für die gesamte Abbildung schreibt man noch (in etwas redundanter Weise)

$$f : X \rightarrow Y \text{ oder } X \xrightarrow{f} Y$$

Man nennt Abbildungen manchmal auch *Funktionen*, insbesondere, wenn ihr Wertebereich ein Zahlbereich ist (bisher haben wir nur von $\mathbb{N}$ und $\mathbb{Z}$ geredet, es ist aber sicher jedem auch $\mathbb{Q}$, $\mathbb{R}$ evtl. sogar $\mathbb{C}$ geläufig, s. Kapitel 2). Der Unterschied zwischen einer Abbildung (oder Funktion) und ihrem Graphen ist formal von psychologischer Natur. Wichtig ist es aber festzuhalten (weil es in der Schule häufig unterschlagen wird), dass

1. die Angabe des Definitionsbereichs zur Definition einer Abbildung dazugehört;
2. Funktionen sowie andere Abbildungen durch Abbildungsvorschriften und im Allgemeinen nicht durch “Funktionsterme” definiert werden.

Beispiel.

$$\begin{aligned} f : \mathbb{R} &\rightarrow \{0, 1\} \\ x &\mapsto f(x) := \begin{cases} 0 & \text{falls } x \text{ rational ist} \\ 1 & \text{falls } x \text{ irrational ist} \end{cases} \end{aligned} \quad (4.1)$$

Hier einige universelle Konventionen über Abbildungen.

Definition 4.2. (i) Wir schreiben $\text{Abb}(X, Y)$ für die Menge aller Abbildungen von X nach Y . ($\text{Abb}(X, Y)$ ist nicht-leer ausser wenn $Y = \emptyset$ und $X \neq \emptyset$.)

(ii) Für jede Menge X heisst die Abbildung $\text{id}_X : X \rightarrow X$, $\text{id}_X(x) := x$ für alle $x \in X$ die *Identität* auf X . (Für $X = \emptyset$ definiert man besser $\text{id}_\emptyset = \emptyset \in \mathcal{P}(\emptyset \times \emptyset)$.)

(iii) Sind X, Y, Z drei Mengen und $f \in \text{Abb}(X, Y)$, $g \in \text{Abb}(Y, Z)$, so definiert $g \circ f(x) := g(f(x))$ eine Abbildung $g \circ f : X \rightarrow Z$, genannt die *Komposition* von f und g . Gesprochen: g nach f .

In diese Definition haben wir, wie es häufig passiert, Aussagen “verpackt”, die man eigentlich beweisen müsste. Im letzten Teil zum Beispiel heisst es, dass die Menge $H \circ G := \{(x, z) \mid \exists y \in Y : (x, y) \in G \wedge (y, z) \in H\}$ die Bedingungen an den Graphen einer Abbildungen von X nach Z im Sinne von Def. 4.1 erfüllt.

Lemma 4.3. *Komposition von Abbildungen ist assoziativ: $h \circ (g \circ f) = (h \circ g) \circ f$.*

Beweis. Ausschreiben! □

Weiterhin ist es von Interesse, inwiefern wir nach “Ersetzen” von x durch $y = f(x)$, x aus y “zurückgewinnen” können.

Definition 4.4. Ist $f : X \rightarrow Y$ eine Abbildung, so heisst $g : Y \rightarrow X$

- (i) *Linksinverse* von f , falls $g \circ f = \text{id}_X$
- (ii) *Rechtsinverse* von f , falls $f \circ g = \text{id}_Y$
- (iii) *Inverse* oder *Umkehrabbildung* von f , falls g sowohl Linksinverse als auch Rechtsinverse von f ist.

Diese Begriffe spielen in interessanter Weise zusammen mit dem Verhalten von Abbildungen auf den Teilmengen des Definitions- und Wertebereichs.

Definition 4.5. Sei $f : X \rightarrow Y$ eine Abbildung.

- (i) Das *Bild* einer Teilmenge $A \subset X$ unter f ist die Menge

$$f(A) := \{y \in Y \mid \exists x \in A : f(x) = y\} \quad (4.2)$$

Das Bild von X (also des gesamten Definitionsbereichs) heisst auch Bild von f .

- (ii) Das *Urbild* einer Teilmenge $B \subset Y$ unter f ist die Menge

$$f^{-1}(B) := \{x \in X \mid f(x) \in B\} \quad (4.3)$$

- (iii) f heisst *injektiv*, falls $\forall x_1, x_2 \in X : f(x_1) = f(x_2) \Rightarrow x_1 = x_2$. Dies ist gleichbedeutend damit, dass für alle $y \in Y$ $f^{-1}(\{y\})$ höchstens ein Element enthält.
- (iv) f heisst *surjektiv*, falls $\forall y \in Y : \exists x \in X : f(x) = y$. Dies ist gleichbedeutend damit, dass für alle $y \in Y$ $f^{-1}(\{y\})$ nicht-leer ist, oder auch, dass $f(X) = Y$.
- (v) f heisst *bijektiv*, falls f injektiv und surjektiv ist.

Proposition 4.6. (i) Eine Abbildung $f : X \rightarrow Y$ ist genau dann bijektiv, wenn eine Umkehrabbildung $g : Y \rightarrow X$ existiert. Die Umkehrabbildung ist bijektiv und eindeutig. Sie wird üblicherweise als $g = f^{-1}$ notiert.

(ii) Ist $X \neq \emptyset$, so ist eine Abbildung $f : X \rightarrow Y$ genau dann injektiv, wenn sie eine Linksinverse besitzt.

(iii) Eine Abbildung $f : X \rightarrow Y$ ist genau dann surjektiv, wenn sie eine Rechtsinverse besitzt.

Die Umkehrabbildung einer bijektiven Abbildung ist zu unterscheiden von den für eine beliebige Abbildung $f : X \rightarrow Y$ durch die Vorschriften (4.2) und (4.3) definierten Abbildungen $f : \mathcal{P}(X) \rightarrow \mathcal{P}(Y)$ und $f^{-1} : \mathcal{P}(Y) \rightarrow \mathcal{P}(X)$. Der Missbrauch der Notation führt selten zu Problemen.

Beweis von Prop. 4.6. (i) “ $\Rightarrow$ ”:⁸ Angenommen, f besitzt eine Umkehrabbildung, g . Ist dann $f(x_1) = f(x_2)$, so folgt durch Anwenden von g auf beiden Seiten, dass $g \circ f(x_1) = g \circ f(x_2)$ und daraus wegen $g \circ f = \text{id}_X$, dass $x_1 = x_2$, also ist f injektiv. Ist dann $y \in Y$ beliebig, so gilt wegen $f \circ g = \text{id}_Y$, dass $f(g(y)) = y$. Es existiert also für jedes y ein x (nämlich $x = g(y)$) so, dass $f(x) = y$. Also ist f surjektiv.

“ $\Leftarrow$ ”: Betrachte für gegebenes f den “umgedrehten Graphen”,

$$G^{-1} := \{(f(x), x) \mid x \in X\} \subset Y \times X \quad (4.4)$$

⁸Der Beweis einer Aussage “ A genau dann, wenn B “, symbolisiert als $B \Leftrightarrow A$, kann in natürlicher Weise in zwei Teile gegliedert werden: Man zeigt zunächst die “direkte Implikation” A wenn B , d.h. $B \Rightarrow A$, und dann die “umgekehrte Implikation” wenn A dann B , d.h. $B \Leftarrow A$.

§ 4. ABBILDUNGEN

Weil f surjektiv ist, existiert für jedes $y \in Y$ ein Element von G^{-1} mit erstem Eintrag y . Der zweite Eintrag ist eindeutig bestimmt, weil f injektiv ist. G^{-1} definiert also eine Abbildung $g : Y \rightarrow X$. Ausserdem ist für jedes $y \in Y$ $g(y)$ genau dasjenige x , für das $f(x) = y$. Das heisst also $f(g(y)) = y$ für alle $y \in Y$. Aus dem gleichen Grund ist für jedes $x \in X$ $g(f(x)) = x$. g ist also eine Umkehrabbildung von f . Ist $\tilde{g}$ eine weitere Umkehrabbildung von f , so folgt aus $y = f(\tilde{g}(y))$ durch Anwenden von g auf beiden Seiten, dass $g(y) = \tilde{g}(y)$ für alle $y \in Y$. Die Umkehrabbildung ist also eindeutig bestimmt.

(ii) “ $\Rightarrow$ ”: Folgt dem ersten Teil des Beweises von (i).

“ $\Leftarrow$ ”: Angenommen f ist injektiv. Wir wählen ein beliebiges $x_* \in X$ und setzen $g(y) := x_*$ für alle $y \in Y \setminus f(X)$. Für jedes der anderen $y \in f(X)$ existiert genau ein $x \in X$ so, dass $f(x) = y$. Wir setzen $g(y) := x$ für solche y . Dann ist g wohldefiniert und es gilt $g(f(x)) = x$ für alle $x \in X$.

(iii) Die Idee des Beweises ist genau wie bei (ii). Die Konstruktion der Rechtsinversen hängt allerdings ab vom sog. “Auswahlaxiom”, auf das wir gleich kurz zu sprechen kommen. $\square$

Beispiel. Die Abbildung $h : \mathbb{N} \rightarrow \mathbb{N}$ $h(n) := n + 1$ ist injektiv. Die Funktion $j : \mathbb{N} \rightarrow \mathbb{N}$, definiert durch $j(n) = n - 1$ falls $n > 1$ und $j(1) = 1$ ist Linksinverse von h , jedoch nicht Rechtsinverse.

Die direkten Implikationen der Proposition 4.6 lassen sich auch gebündelt einfacher formulieren.

Lemma 4.7. *Sind $f : X \rightarrow Y$ und $g : Y \rightarrow X$ zwei Abbildungen mit der Eigenschaft, dass $g \circ f = \text{id}_X$, dann ist f injektiv und g surjektiv.*

Beweis. Aus $f(x_1) = f(x_2)$ folgt durch Anwenden von g auf beiden Seiten, dass $x_1 = x_2$. Also ist f injektiv. Für jedes $x \in X$ hat $y = f(x) \in Y$ die Eigenschaft, dass $g(y) = x$. Also ist g surjektiv. $\square$

Mit diesen Einsichten gerüstet reichen wir noch einige mengentheoretische Begriffe nach. Achtung: Wir begeben uns zwischendurch bis an die Schwelle grösserer Komplikationen, unter anderem um zu rechtfertigen, warum uns dies nicht kümmern muss. Zunächst erläutern wir den Zusammenhang zwischen Abbildungen und Äquivalenzrelationen.

Proposition 4.8. *Sei X eine Menge mit Äquivalenzrelation $\sim$ (s. Def. 3.5), und $\mathcal{C}$ die Menge der Äquivalenzklassen bezüglich $\sim$.*

(i) *Für alle $C_1, C_2 \in \mathcal{C}$ gilt entweder $C_1 = C_2$ oder $C_1 \cap C_2 = \emptyset$.*

(ii) *Für jedes $x \in X$ existiert ein eindeutiges $C_x \in \mathcal{C}$ so, dass $x \in C_x$.*

(iii) *Die Abbildung $X \rightarrow \mathcal{C}$, $x \mapsto C_x$ ist surjektiv. Sie heisst kanonische Projektion.*

(iv) *Jede Abbildung $p : Y \rightarrow Z$ definiert via $y_1 \sim_p y_2 \Leftrightarrow p(y_1) = p(y_2)$ eine Äquivalenzrelation auf Y . Ist p surjektiv, so ist die Abbildung $p^{-1} : Z \rightarrow Y/\sim_p$, $z \mapsto p^{-1}(\{z\})$ eine (kanonische) Bijektion von Z auf die Menge der Äquivalenzklassen bezüglich $\sim_p$.*

Man sagt “ X zerfällt in disjunkte Äquivalenzklassen”. Die Menge der Äquivalenzklassen heisst auch *Faktormenge bezüglich $\sim$* (oder Quotient von X nach $\sim$) und wird als $X/\sim$ geschrieben. Man nennt C_x auch *Äquivalenzklasse von x* und schreibt hierfür normalerweise $[x]$. Dies bedeutet aber nicht, dass x ein in irgendeiner Weise ausgezeichneter Repräsentant seiner Äquivalenzklasse ist!!!

Beispiel. Für $n \in \mathbb{N}$ und $k \in \mathbb{Z}$ schreibt man für die Restklasse von k modulo n häufig $[k] = k + n\mathbb{Z}$. Durch diese Notation wird aber *nicht* impliziert, dass k zwischen 1 und n liegt! Für die Menge der Restklassen schreiben wir $\mathbb{Z}/n\mathbb{Z}$, manchmal auch $\mathbb{Z}/n$ oder sogar $\mathbb{Z}_n$.

Beweis von Prop. 4.8. (i) Angenommen $C_1 \cap C_2 \neq \emptyset$. Dann existiert ein $x_* \in C_1 \cap C_2$. Ist dann $x \in C_1$, so gilt wegen 3.5 (ii), dass $x \sim x_*$. Da $x_* \in C_2$ folgt aus 3.5 (iii), dass auch $x \in C_2$. Umgekehrt geht genauso, d.h. $C_1 = C_2$.

(ii) Wir setzen $C_x = \{\tilde{x} \in C \mid \tilde{x} \sim x\}$ und überprüfen, dass C_x eine Äquivalenzklasse mit $x \in C_x$ ist.

(iii) Wohldefiniertheit folgt aus (ii), Surjektivität aus der Tatsache, dass jede Äquivalenzklasse nicht leer ist.

(iv) Die Reflexivität, Symmetrie und Transitivität von $\sim_p$ folgen aus denen von $=$. Bijektivität von p^{-1} sei Übungsaufgabe. $\square$

Nun erweitern wir etwas die mengentheoretischen Operationen von Def. 2.4.

Definition 4.9. Sind $\mathcal{M}$ und I zwei beliebige Mengen, so heisst eine Abbildung $M_- : I \rightarrow \mathcal{M}$, geschrieben als $I \ni i \mapsto M_i \in \mathcal{M}$ auch *I -indizierte Familie von Objekten der Sorte $\mathcal{M}$* , und wird dann als $(M_i)_{i \in I}$ notiert. Ist $\mathcal{M}$ speziell eine Menge von Mengen,⁹ so definiert man

- (i) die Vereinigung der Familie als $\bigcup_{i \in I} M_i := \{x \mid \exists i : x \in M_i\}$
- (ii) den Durchschnitt der Familie als $\bigcap_{i \in I} M_i := \{x \mid \forall i : x \in M_i\}$
- (iii) das kartesische Produkt $\times_{i \in I} M_i := \{(m_i)_{i \in I} \mid \forall i : m_i \in M_i\}$

Beispiel. In der Situation von Prop. 4.8 schreiben wir etwa $X = \bigcup_{C \in \mathcal{C}} C$.

Bemerkungen. Die Vereinigung beliebiger Mengen bietet keinerlei Schwierigkeiten. Ist $I = \emptyset$, so ist $\bigcup_{i \in \emptyset} M_i = \emptyset$. Hingegen ist der Durchschnitt einer durch die leere Menge indizierte Familie im Allgemeinen nicht definiert, oder allenfalls gleich dem gesamten Diskursuniversum. Beim kartesischen Produkt gibt es ein ganz unerwartetes Problem, wenn I “sehr gross” wird, weil man dann im Allgemeinen I -indizierte Familien von Elementen $(m_i \in M_i)_{i \in I}$ nicht *explizit konstruieren* kann. Die nicht-triviale Aussage, “das kartesische Produkt beliebiger Familien nicht-leerer Mengen ist nicht-leer”, wird als *Auswahlaxiom* bezeichnet. Wir werden es nur selten benutzen, aber sonst nicht in Frage stellen. Der wesentliche Grund hierfür ist, dass wir *für die Physik* keine beliebig grossen Mengen brauchen.

Definition 4.10. Sei X eine Menge.

- (i) X heisst *gleichmächtig* zu einer Menge Y , geschrieben $|X| = |Y|$ oder $\#X = \#Y$, wenn eine bijektive Abbildung $f : X \rightarrow Y$ existiert.

⁹Prinzipiell lässt sich jedes mathematische Objekt als Menge begreifen. Insofern handelt es sich hier vor allem um eine psychologische Standortbestimmung.

§ 4. ABBILDUNGEN

- (ii) X heisst *endlich*, wenn sie gleichmächtig zu $\emptyset$ oder zu $\{1, 2, \dots, n\}$ für ein $n \in \mathbb{N}$ ist. Man schreibt in diesen Fällen $|X| = \#X = n$, bzw. $|\emptyset| = 0$.
- (iii) X heisst *abzählbar*, wenn sie entweder endlich ist oder gleichmächtig zu $\mathbb{N}$.

Die Schreibweise ist dadurch gerechtfertigt, dass Gleichmächtigkeit eine Äquivalenzrelation ist. Man sagt auch, X ist *höchstens gleichmächtig* zu Y , wenn eine injektive Abbildung $X \rightarrow Y$ existiert, und schreibt dann $|X| \leq |Y|$. Die Eigenschaften dieser Relation sind allerdings subtil (so ist z.B. die Aussage, dass $|X| \leq |Y| \wedge |Y| \leq |X| \Rightarrow |X| = |Y|$, Inhalt des berühmten Bernstein-Schröder Theorems). Viel einfacher und enorm nützlich ist die folgende Aussage.

Proposition 4.11. *Seien X und Y endliche Mengen mit $\#X = \#Y$, und $f : X \rightarrow Y$ eine Abbildung. Dann sind äquivalent:*

- (i) f ist injektiv
- (ii) f ist surjektiv
- (iii) f ist bijektiv

Wir beleuchten zunächst einen Spezialfall.

Lemma 4.12. *Sei X eine endliche Menge.*

- (i) *Ist $A \subsetneq X$ eine echte Teilmenge, so gilt $\#A < \#X$.*
- (ii) *Ist $X = \bigcup_{i \in I} X_i$ eine disjunkte Zerlegung von X , so gilt*

$$\#X = \sum_{i \in I} \#X_i \quad (4.5)$$

Beweis. Durch elementare Einsicht in die natürlichen Zahlen: Steht X in Bijektion zu $\{1, 2, \dots, n\}$, und ist $A \subsetneq X$, so lässt sich durch Abzählen eine Bijektion angeben von A zu $\{1, 2, \dots, k\}$ für ein $k < n$. Daraus folgt die Behauptung (i). Ebenso induziert eine disjunkte Zerlegung von X eine disjunkte Zerlegung von $\{1, 2, \dots, n\}$, für die die Formel (ii) aus den Gesetzen der Arithmetik folgt. $\square$

Beweis von Prop. 4.11. Nach Def. 4.5 gilt (iii) $\Rightarrow$ (i) und (iii) $\Rightarrow$ (ii). Es genügt also zu zeigen, dass Injektivität äquivalent zu Surjektivität ist. Angenommen, f ist injektiv. Dann ist f eine Bijektion von X auf $f(X)$ (Übungsaufgabe) und daher gilt $\#f(X) = \#X = \#Y$. Aus dem Lemma folgt $f(X) = Y$, d.h. f ist surjektiv. Ist umgekehrt f surjektiv, dann gilt $\#f^{-1}(\{y\}) \geq 1$ für alle $y \in Y$. Aus Prop. 4.8 (iv) folgt $X = \bigcup_{y \in Y} f^{-1}(\{y\})$ und damit aus dem Lemma $\#f^{-1}(\{y\}) = 1$ für alle $y \in Y$, d.h. f ist injektiv. $\square$

Siehe S. 8 für ein interessantes Beispiel. Zum Abschluss teilen wir ein paar Einsichten zur Potenzmenge.

Theorem 4.13. *Sei X eine Menge.*

- (i) *Die Abbildung $\mathcal{P}(X) \rightarrow \text{Abb}(X, \mathbb{Z}/2\mathbb{Z})$, gegeben durch*

$$\mathcal{P}(X) \ni A \mapsto \chi_A \in \text{Abb}(X, \mathbb{Z}/2), \quad \chi_A(x) := \begin{cases} [1] & \text{falls } x \in A \\ [0] & \text{falls } x \notin A \end{cases} \quad (4.6)$$

ist bijektiv.

(ii) $\mathcal{P}(X)$ ist endlich genau dann wenn X endlich ist. Dann gilt $\#\mathcal{P}(X) = 2^{\#X}$

(iii) Im Allgemeinen gilt stets $|X| \leq |\mathcal{P}(X)|$ aber $|\mathcal{P}(X)| \neq |X|$.

Beweis. wird in der Plenarübung vorgeführt □

Motiviert durch Theorem 4.13 (ii) denkt man sich die Spezialisierung des kartesischen Produktes in 4.9 (iii) auf den Fall, dass alle M_i “gleich Y ” sind, und $I = X$, als

$$Y^X = \bigtimes_{x \in X} Y_x = \{(y_x)_{x \in X}\} = \{f : X \rightarrow Y\} = \text{Abb}(X, Y) \quad (4.7)$$

Hierfür bedarf es nicht des Auswahlaxioms.

KAPITEL 2

ZAHLEN

Die Essenz des ersten Kapitels ist, dass wir die natürlichen Zahlen begreifen können als Äquivalenzklassen von nicht-leeren endlichen¹⁰ Mengen unter der Äquivalenzrelation, dass $A \sim B$, falls eine Bijektion von A nach B existiert.¹¹ Als Kinder haben wir ausserdem eingesehen, dass die disjunkte Vereinigung endlicher Mengen der *Addition*, und das kartesische Produkt der *Multiplikation* natürlicher Zahlen entspricht.¹² Die Inklusion von Mengen induziert die *totale Anordnung* der natürlichen Zahlen und es gelten die bekannten Rechenregeln.

Addition und Multiplikation mit einer konstanten Zahl sind Abbildungen von $\mathbb{N}$ auf sich selbst, die injektiv, jedoch im Allgemeinen nicht surjektiv sind. Der Wunsch der *Buchhaltung*, dass bei derartigen Operationen keine Information verloren geht, hat uns später dazu geführt, zunächst zu den ganzen Zahlen, und dann zu den rationalen Zahlen überzugehen. Beide lassen sich “elegant” als Faktormenge schreiben.¹³

$$\begin{aligned} \mathbb{Z} &:= \mathbb{N} \times \mathbb{N} / \sim_+ \text{ wobei } (a_1, b_1) \sim_+ (a_2, b_2) \Leftrightarrow a_1 + b_2 = a_2 + b_1 \\ &\simeq \{ \{ (a + n, n) \mid n \in \mathbb{N} \} \mid a \in \mathbb{N} \} \sqcup \{ \{ (n, n) \mid n \in \mathbb{N} \} \} \\ &\quad \sqcup \{ \{ (n, a + n) \mid n \in \mathbb{N} \} \mid a \in \mathbb{N} \} \end{aligned} \quad (4.1)$$

$$\mathbb{Q} := \mathbb{Z} \times (\mathbb{Z} \setminus \{0\}) / \sim. \text{ wobei } (p_1, q_1) \sim (p_2, q_2) \Leftrightarrow p_1 q_2 = p_2 q_1 \quad (4.2)$$

$\mathbb{N}$ ist in kanonischer Weise in $\mathbb{Z}$, und $\mathbb{Z}$ in kanonischer Weise in $\mathbb{Q}$ enthalten, und die Rechenoperationen lassen sich auf die nächstgrössere Menge *fortsetzen*¹⁴ unter Einhaltung der Rechenregeln (Kommutativität, Assoziativität, Distributivgesetz).

¹⁰Diese Charakterisierung ist natürlich zirkulär, da wir Endlichkeit von Mengen genau mit Hilfe der natürlichen Zahlen definiert haben. Dies lässt sich nur axiomatisch (Peano) oder operationell (Vorzeigen) durchbrechen.

¹¹ Diese Identifikation funktioniert etwas natürlicher für $\mathbb{N}_0$ als für $\mathbb{N}$. So ist $0 = \emptyset$, $1 = \{0\} = \{\emptyset\}$, $2 = \{0, 1\} = \{\emptyset, \{\emptyset\}\}$, $3 = \{0, 1, 2\} = \{\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\}\} \dots$

¹²Wir hatten unter Def. 2.4 festgehalten, dass das kartesische Produkt nicht kommutativ ist, da im Allgemeinen $A \times B \neq B \times A$. Allerdings gibt es aber stets eine (kanonische) Bijektion, die $(a, b) \in A \times B$ auf $(b, a) \in B \times A$ abbildet. Das Produkt auf $\mathbb{N}$ ist also natürlich kommutativ.

¹³Das Symbol $\simeq$ bedeutet, dass zwischen der linken und der rechten Seite eine Bijektion existiert, die mit allen weiteren Vereinbarungen (hier: Rechenregeln) verträglich ist.

¹⁴ Ist $f : X \rightarrow Y$ eine Abbildung und $U \subset X$, so heisst die Abbildung $f|_U : U \rightarrow Y$, definiert durch $f|_U(u) := f(u)$ die *Einschränkung* von f auf U . Ist $X \subset Z$, so heisst eine Abbildung $\tilde{f} : Z \rightarrow Y$ *Fortsetzung* von f nach Z , falls $\tilde{f}|_X = f$. Die Einschränkung ist immer wohldefiniert und eindeutig und wird häufig mit dem gleichen Symvol bezeichnet. Eine Fortsetzung existiert falls $Y \neq \emptyset$, ist ohne weitere Bedingungen nicht eindeutig, und wird deshalb normalerweise mit einem anderen Symvol notiert. Der Zusammenhang zwischen Rechenoperationen und Abbildungen wird sofort erklärt.

Die *Algebra* untersucht *Verallgemeinerungen* solcher Rechenregeln und die Lösung von Problemen in “endlich vielen Schritten”.

Für die Zwecke der exakten Naturwissenschaften ist es essentiell, dass sich die Anordnung von $\mathbb{N}$ zu einem *Grössenvergleich* auf $\mathbb{Q}$ fortsetzen lässt, der mit den Rechenoperationen verträglich ist. Dabei stellt man fest:

1. Für alle praktischen Zwecke reicht es aus, physikalische Grössen *approximativ* zu messen. Im Allgemeinen bestehen solche Grössen aus *mehreren* Zahlen.
2. Man kann die Messunsicherheit (mit ausreichend €€) anscheinend beliebig *verkleinern*.
3. Die zugehörigen mathematischen Probleme lassen sich in *endlich vielen* Schritten *nicht exakt* lösen.

Diese Beobachtungen dienen als Motivation zur Einführung erst der *reellen* und anschliessend der *komplexen Zahlen* als Grundlage der *Analysis*, und können auch zur Begründung *linearer Strukturen* benutzt werden. Wir schliessen dieses Kapitel mit einem kurzen Einblick in diese Zusammenhänge, fangen aber zunächst klein an.

§ 5 Gruppen, Ringe, Körper

Definition 5.1. (i) Eine (binäre, innere) *Verknüpfung* auf einer Menge A ist eine Abbildung $f : A \times A \rightarrow A$, geschrieben $(a_1, a_2) \mapsto f(a_1, a_2)$.

(ii) Eine Verknüpfung f heisst *assoziativ* falls $f(f(a_1, a_2), a_3) = f(a_1, f(a_2, a_3))$.

(iii) Eine Menge H zusammen mit einer assoziativen Verknüpfung f heisst *Halbgruppe*.

(iv) Eine innere Verknüpfung / Halbgruppe heisst *kommutativ* oder *abelsch* falls $f(a_1, a_2) = f(a_2, a_1)$ für alle $a_1, a_2 \in H$.

Wir benutzen als Symbol für f unter anderem $\star, +, \times, \cup, \cdot, \circ, \otimes, \dots$ und notieren es statt polnisch normalerweise als Infix: $a_1 \star a_2 := \star(a_1, a_2)$. Die Verwendung von $+$ o.ä. impliziert immer Kommutativität, bei $\times$ oder $\cdot$ ist es manchmal so, manchmal so. Ein Multiplikationssymbol wird häufig auch ganz weggelassen. Im assoziativ-kommutativen Fall wird die Verknüpfung von endlich(!) vielen Elementen auch geschrieben als

$$\begin{aligned} \star_{i=1}^n a_i &:= a_1 \star a_2 \star \cdots \star a_n \\ \text{speziell für } f = + \text{ auch: } \sum_{i=1}^n a_i &:= a_1 + a_2 + \cdots + a_n \\ \text{und für } f = \cdot \text{ oder } \times: \prod_{i=1}^n a_i &:= a_1 \cdot a_2 \cdots a_n \end{aligned} \tag{5.1}$$

Summen und Produkten von unendlich vielen Elementen geben wir (im aktuellen Setting) *keinen Sinn*. Im multiplikativen (kommutativen oder nicht-kommutativen) Fall definiert man für jedes $a \in H$ und $n \in \mathbb{N}$

$$a^{\star n} = a^n := \underbrace{a \star a \star \cdots \star a}_{n \text{ mal}}, \text{ es gilt } a^n \star a^m = a^{n+m} \tag{5.2}$$

§ 5. GRUPPEN, RINGE, KÖRPER

Aus $a_1 \star a_2 = a_2 \star a_1$ folgt $(a_1 \star a_2)^n = a_1^n \star a_2^n$.

Beispiele. $(\mathbb{N}, +)$, $(\mathbb{Z}, \cdot)$, $(\text{Abb}(X, X), \circ)$, ... sind alles Halbgruppen. Die beiden ersten sind kommutativ.

Definition 5.2. Sei $(H, \star)$ eine Halbgruppe.

(i) Ein Element $e \in H$ heisst *neutrales Element* (oder *Eins* bezüglich $\star$, wir schreiben auch 1_H), falls $\forall a \in H : e \star a = a = a \star e$. Ein neutrales Element (falls es existiert) ist notwendigerweise eindeutig.

(ii) Ist e neutrales Element und $a, b \in H$ mit $a \star b = e$, so heisst a *Links inverses* von b und b *Rechts inverses* von a .

(iii) Ist e neutrales Element und besitzt $a \in H$ ein Links inverses und ein Rechts inverses, so heisst a *invertierbar*. Links- und Rechtsinverse sind dabei gleich und werden als a^{-1} notiert. Es gilt $(a^{-1})^{-1} = a$.

Im abelsch-additiven Fall notieren wir das neutrale Element (falls es existiert) als ' 0_H ' und das Inverse (falls es existiert) als ' $-a$ '. Ausserdem vereinbaren wir natürlich $b - a := b + (-a)$.

Beweis der inkludierten Behauptungen. (i) Ist $\tilde{e}$ ein anderes neutrales Element, so folgt $\tilde{e} = \tilde{e} \star e = e$. (iii) Ist $(a \star b = e) \wedge (c \star a = e)$, so folgt

$$b = e \star b = \overbrace{c \star a}^e \star b = c \star e = c \quad (5.3)$$

$(a^{-1})^{-1} = a$ ist gleichbedeutend mit $a \star a^{-1} = a^{-1} \star a = e$. Das ist aber genau die Bedingung dafür, dass a^{-1} Inverses von a ist. $\square$

Warnung: Allein aus der *Existenz* eines Linksinversen folgt noch nicht die Existenz eines Rechtsinversen! (Man könnte auch noch zwischen links- und rechtsneutralen Elementen unterscheiden, das wird aber etwas unübersichtlich.)

Definition 5.3. Eine *Gruppe* $(G, \star)$ ist eine (nicht-leere) Halbgruppe mit neutralem Element $e \in G$, in der jedes Element invertierbar ist. Eine Gruppe heisst *abelsch*, falls $\star$ kommutativ ist. Wir adressieren eine solche Gruppe auch häufig nur mit G .

Beispiele 5.4. 1. $\mathbb{Z} = (\mathbb{Z}, +)$ (mit $e = 0$) und $\mathbb{Q}^\times = (\mathbb{Q} \setminus \{0\}, \cdot)$ (mit $e = 1$) sind abelsche Gruppen.

2. Ist $(H, \star)$ eine Halbgruppe mit neutralem Element $e \in H$, so ist $H^\times := \{a \in H \mid a \text{ invertierbar}\}$ mit (der Einschränkung von) $\star$ eine Gruppe, genannt *die Einheitsgruppe* von H .

3. Ist X eine beliebige Menge, so ist die Menge der *Bijektionen* $\text{Bij}(X) = \{f \in \text{Abb}(X, X) \mid f \text{ invertierbar}\}$ mit Komposition von Abbildungen als Verknüpfung eine Gruppe mit neutralem Element $e = \text{id}_X$.

In den Nachweis von 2. geht ein, dass die Verknüpfung $a_1 \star \dots \star a_n$ von invertierbaren Elementen wieder invertierbar ist mit Inversen $a_n^{-1} \star \dots \star a_1^{-1}$. 3. ist ein Spezialfall von 2. Die Definition (5.2) setzen wir mit $a^0 = e$ fort zu a^n für alle $n \in \mathbb{Z}$.

Definition 5.5. Ist $X \simeq \{1, 2, \dots, n\}$ eine (nicht-leere) endliche Menge, so heisst $\text{Bij}(X)$ die *symmetrische Gruppe vom Grad n* , geschrieben: S_n .

Während die Bedeutung von (nicht-abelschen!) Gruppen für die Physik gar nicht genug betont werden kann (s. PTP2–4, HöMa3, etc.), interessieren wir uns aus den eingangs genannten Gründen zunächst für Strukturen mit zwei Verknüpfungen, die die Rollen von Addition und Multiplikation spielen sollen.

Definition 5.6. Ein Tupel $(R, +, \star)$ heisst *Ring*, wenn gilt

- (i) $(R, +)$ ist eine abelsche Gruppe;
- (ii) $(R, \star)$ ist eine Halbgruppe;
- (iii) Es gelten die Distributivgesetze: Für alle a_1, a_2, a_3 :

$$a_1 \star (a_2 + a_3) = a_1 \star a_2 + a_1 \star a_3, \quad (a_1 + a_2) \star a_3 = a_1 \star a_3 + a_2 \star a_3 \quad (5.4)$$

Das neutrale Element der Addition wird normalerweise mit 0_R notiert. Die Bezeichnung *Ring mit Eins* (man sagt auch “unitärer Ring”; das neutrale Element wird mit 1_R notiert) und die Begriffe Kommutativität und Invertierbarkeit beziehen sich auf die multiplikative Verknüpfung $\star$. Beachte: Die Einheitengruppe $R^\times \subset R$ ist im Allgemeinen kein Ring(!).

Beispiel. Ist $n \in \mathbb{N}$, so wird die Menge $R = \mathbb{Z}/n\mathbb{Z}$ der Restklassen modulo n (s. Beispiel 3.6) mit

$$\begin{aligned} [n_1] + [n_2] &:= [n_1 + n_2] \\ [n_1] \cdot [n_2] &:= [n_1 \cdot n_2] \end{aligned} \quad (5.5)$$

und $0_R = [0]$, $1_R = [1]$ zu einem kommutativen Ring mit Eins. Man überprüfe Wohldefiniertheit und Rechengesetze.

Definition 5.7. Ein Tupel $(K, +, \cdot)$ heisst *Körper*, wenn gilt

- (i) $(K, +, \cdot)$ ist ein kommutativer Ring mit Eins 1_K .
- (ii) Alle Elemente ausser der 0_K sind invertierbar, d.h. $K^\times = K \setminus \{0_K\}$.

Bemerkung. In jedem Ring gilt $\forall a \in R : 0_R \star a + 0_R \star a = (0_R + 0_R) \star a = 0_R \star a = 0_R \star a + 0_R$, d.h. $0_R \star a = 0_R$, und ebenso $a \star 0_R = 0_R$, insbesondere gilt $0_R \star 0_R = 0_R$. Daraus folgt: 0_R ist genau dann invertierbar, wenn $0_R = 1_R$ und in diesem Fall gilt $a = a \star 1_R = a \star 0_R = 0_R$ für alle a , d.h. R besteht nur aus einem Element. In einem Körper ist die 0_K nicht invertierbar, und daher gilt auch $0_K \neq 1_K$. Ein Körper besteht aus mindestens zwei Elementen.

Beispiele. $(\mathbb{Q}, +, \cdot)$ mit $0_{\mathbb{Q}} = 0$, $1_{\mathbb{Q}} = 1$ ist ein Körper. Der oben definierte Ring $\mathbb{F}_n := (\mathbb{Z}/n\mathbb{Z}, +, \cdot)$ ist genau dann ein Körper, wenn n eine Primzahl ist.

Bottom line: Wir können in einem Körper rechnen wie mit den rationalen Zahlen aus der Schule gewohnt. In den folgenden §§ beschäftigen wir uns vor allem mit der Erweiterung von $\mathbb{Q}$ zu $\mathbb{R}$ und $\mathbb{C}$. Es gibt “in der Natur” aber noch sehr viel mehr (interessante!) Körper, die für die Physik allerdings weniger wichtig sind.¹⁵

¹⁵Ein Körper heisst *von endlicher Charakteristik*, wenn es eine natürliche Zahl n gibt mit der Eigenschaft, dass $\underbrace{1_K + 1_K + \dots + 1_K}_{n\text{-mal}} = 0_K$. In diesem Fall ist die kleinste solche Zahl notwendiger-

weise eine Primzahl, p . Sie heisst die *Charakteristik* von K . Beispiel: $\mathbb{Z}/p\mathbb{Z}$ ist ein endlicher Körper von Charakteristik p . Körper ohne solches n , wie insbesondere $\mathbb{Q}$, $\mathbb{R}$, $\mathbb{C}$, heissen *von Charakteristik 0*.

§ 5. GRUPPEN, RINGE, KÖRPER

Gruppen und (in etwas abgeschwächter oder weniger expliziten Form) allgemeinere Ringe spielen allerdings eine eminente Rolle.

Beispiel 5.8. Ist $(K, +, \cdot)$ ein Körper und $x \notin K$ ein beliebiges sonstiges Symbol, so wird die Menge der “formalen Ausdrücke in der Unbekannten x ”

$$R := \left\{ p = a_n x^n + a_{n-1} x^{n-1} + \cdots + a_1 x + a_0 \mid \right. \\ \left. n \in \mathbb{N}_0 \text{ und } a_i \in K \text{ für alle } i = 0, 1, \dots, n \right\} / \sim \quad (5.6)$$

wobei $a_n x^n + \cdots + a_0 \sim b_m x^m + \cdots + b_0$, falls sich die beiden Ausdrücke nur um Terme der Form $0_K x^k$ unterscheiden, zu einem kommutativen Ring mit Eins, indem man definiert

$$(a_n x^n + a_{n-1} x^{n-1} + \cdots + a_1 x + a_0) + (b_n x^n + b_{n-1} x^{n-1} + \cdots + b_1 x + b_0) \\ := (a_n + b_n) x^n + (a_{n-1} + b_{n-1}) x^{n-1} + \cdots + (a_1 + b_1) x + (a_0 + b_0) \quad (5.7)$$

$$(a_n x^n + a_{n-1} x^{n-1} + \cdots + a_1 x + a_0) \cdot (b_n x^n + b_{n-1} x^{n-1} + \cdots + b_1 x + b_0) \\ := a_n b_n x^{2n} + (a_n b_{n-1} + a_{n-1} b_n) x^{2n-1} + \cdots + (a_1 b_0 + a_0 b_1) x + a_0 b_0, \quad (5.8)$$

Hierbei ist

$$0_R = 0_K \quad (5.9)$$

$$1_R = 1_K \quad (5.10)$$

Man nennt diesen Ring “Polynomring über K ” und schreibt $R = K[x]$. Polynome dienen unter anderem zur Definition von Funktionen als Abbildungen von K nach K . Wir schreiben in dieser Interpretation auch $p(x)$ statt p .

$$p : K \rightarrow K, \quad x \mapsto p(x) = a_n x^n + a_{n-1} x^{n-1} + \cdots + a_1 x + a_0 \quad (5.11)$$

Dies ist wohldefiniert, da die Funktion sich unter Addition von Termen der Form $0_K x^k$ nicht ändert. Der grösste Index d mit der Eigenschaft, dass $a_d \neq 0$, heisst der *Grad* des Polynoms, geschrieben $\deg(p)$. Existiert kein solches a_d (d.h., ist $p = 0_R$), so ist (per Definition) $\deg(p = 0_R) := -\infty$.¹⁶

Definition 5.9. Sei K ein Körper und $p \in K[x]$ ein Polynom. Dann heisst $x_0 \in K$ *Nullstelle* von p , falls $p(x_0) = 0$.

Proposition 5.10. Sei $p(x) \in K[x]$ ein Polynom vom Grad $\deg(p) = d$, $d > 0$, und $x_0 \in K$ eine Nullstelle von p . Dann existiert ein Polynom $p' \in K[x]$ mit Grad $\deg(p') = d - 1$ so, dass $p(x) = (x - x_0)p'(x)$. Insbesondere hat p höchstens d (verschiedene) Nullstellen.

Beweis. Siehe Übungen □

· Die Bezeichnung “Körper” geht auf R. Dedekind (1831-1916) zurück, der englische Begriff “field” auf E. H. Moore (1862-1932).

Would have been smart to present matrices here as another example of things you can multiply.

¹⁶Dies stellt insbesondere sicher, dass $\deg(p \cdot q) = \deg(p) + \deg(q)$ für alle $p, q \in R$.

§ 6 Die reellen Zahlen

Wir bezeichnen ab jetzt die neutralen Elemente eines Körpers einfach mit 0 und 1, lassen bei der Multiplikation den Malpunkt weg, und schreiben Multiplikation mit a^{-1} als Division durch a . Wesentliche Konsequenz der im letzten § entwickelten Axiome ist, dass die Gleichungen

$$a + x = 0, \quad b \cdot y = 1 \quad (6.1)$$

für alle $a \in K$ und $b \in K \setminus \{0\}$ eindeutige Lösungen $x = -a$ und $y = b^{-1} \in K$ besitzen. Dies hatten wir in der Einleitung ja auch als Motivation für die Erweiterung von $\mathbb{N}$ zu $\mathbb{Z}$ und $\mathbb{Q}$ genannt. Der aus rein algebraischer Sicht ebenfalls mögliche Übergang zu Restklassen (s. Fussnote 15) ist aus Sicht der Physik unter anderem deshalb unbrauchbar, weil er nicht mit dem Wunsch nach einer Anordnung des Rechenbereichs zum *quantitativen* Vergleich physikalischer Grössen verträglich ist.

Definition 6.1. Ein *angeordneter Körper* ist ein Körper K zusammen mit einer ausgezeichneten Teilmenge P (der Menge der *positiven* Elemente) derart, dass

$$\begin{aligned} K &= P \cup \{0\} \cup (-P) \\ P + P &\subset P \\ P \cdot P &\subset P \end{aligned} \quad (6.2)$$

Hierbei ist $-P := \{-a \mid a \in P\}$ und gemäss Def. 2.4 bedeutet $\cup$ “disjunkte Vereinigung”, das heisst insbesondere $0 \notin P$ und $P \cap (-P) = \emptyset$. Wir schreiben auch $a > 0$ oder $0 < a$ für $a \in P$ und $a < 0$ oder $0 > a$ für $a \in -P$. Wir zeigen im Folgenden zunächst, dass Begriff und Notation verträglich sind mit der vertrauten Anordnung der rationalen (und reellen) Zahlen auf der Zahlengeraden. Danach verweisen wir kurz auf die Grenzen solcher Anordnungen (das sog. archimedische Axiom). Zum Schluss beschreiben wir den Zusammenhang zwischen Anordnung und eigentlicher Grössenabschätzung.

Lemma 6.2. (i) Durch die Vorschrift $a > b : \Leftrightarrow a - b \in P$ wird auf einem angeordneten Körper eine strenge Totalordnung (im Sinne von Def. 3.3) definiert (transitiv und trichotomisch), bzw. durch $a \geq b : \Leftrightarrow (a > b \text{ oder } a = b)$ eine Totalordnung (anti-symmetrisch, transitiv und total) und es gelten die üblichen (zum Teil empfindlichen) Rechenregeln für Ungleichungen, insbesondere

$$\begin{aligned} \text{(ii) Aus } a > b \text{ folgen } &\begin{cases} \frac{1}{a} < \frac{1}{b} & \text{falls } b > 0 \\ a + c > b + c & \forall c \in K \\ ac \geq bc & \text{je nach dem } c \geq 0 \end{cases} \\ \text{(iii) Aus } a > b \text{ und } c > d \text{ folgen } &\begin{cases} a + c > b + d & \text{in jedem Fall} \\ ac > bd & \text{falls } b > 0 \text{ und } d > 0 \end{cases} \\ \text{(iv) Für } a \neq 0 \text{ gilt } &a^2 > 0. \text{ Insbesondere ist} \end{aligned}$$

$$1 > 0 \quad (6.3)$$

§6. DIE REELLEN ZAHLEN

Beweis. (i) Die Trichotomie: $\forall a, b \in K$ gilt entweder $a > b$ oder $a = b$ oder $b > a$ folgt aus der dreiteiligen Zerlegung in (6.2). Die Transitivität sehen wir so:

$$\begin{aligned} a > b \text{ und } b > c &\Leftrightarrow a - b \in P \text{ und } b - c \in P \Rightarrow \\ &(\text{wegen } P + P \subset P) \Rightarrow (a - b) + (b - c) = a - c \in P \Leftrightarrow a > c \quad (6.4) \end{aligned}$$

(ii) Wäre $\frac{1}{a} - \frac{1}{b} > 0$ so folgte durch Multiplikation mit der positiven Zahl ab , dass $b > a$. Aus $a - b > 0$ folgt $a + c - b - c > 0$ und daraus per Definition $a + c > b + c$. Die letzte Zeile folgt aus $a - b > 0$ durch Multiplikation mit c bzw. $-c$.

(iii) $a - b > 0$ und $c - d > 0$ implizieren $a + c - b - d > 0$. Aus (ii) folgt zunächst $ac > bc$, und anschliessend $bc > bd$, zusammen $ac > bd$.

(iv) Falls $a > 0$ so folgt $a^2 > 0$ direkt aus $P \cdot P \subset P$. Falls $a < 0$ so ist per Definition $-a > 0$. Wegen $(-1) \cdot a + a = (-1 + 1) \cdot a = 0$ gilt $-a = (-1) \cdot a$ und wegen $(-a)^2 - a^2 = (-a)^2 + (-a) \cdot a = (-a) \cdot (-a + a) = 0$ gilt $a^2 = (-a)^2 > 0$, wieder wegen $P \cdot P \subset P$. Es gilt also insbesondere $1 = (-1)^2 > 0$ $\square$

Allein aus diesen Eigenschaften, d.h. insbesondere ohne Benutzung des archimedischen Axioms Def. 6.7, folgt die folgende Ungleichung, die als einfaches Hilfsmittel zum Einstieg in viele Grössenabschätzungen dient.

Lemma 6.3 (Bernoullische Ungleichung). *Für alle $x \geq -1$ in einem angeordneten Körper K gilt für alle $n \in \mathbb{N}$:*

$$(1 + x)^n \geq 1 + nx \quad (6.5)$$

Beweis. Auf der rechten Seite ist $(1 + x)^n$ durch Rekursion definiert, und daher beweisen wir die Aussage auch durch vollständige Induktion. Für $n = 1$ ist die Aussage klar, der Schluss von n auf $n + 1$ ergibt sich wegen $1 + x \geq 0$ so:

$$(1 + x)^{n+1} \geq (1 + nx)(1 + x) = 1 + (n + 1)x + nx^2 \geq 1 + (n + 1)x \quad (6.6)$$

$\square$

Die Auszeichnung der positiven Elemente ist also gleichbedeutend mit der vertrauten Anordnung der Zahlengeraden. In der Physik ist dies von fundamentaler Bedeutung zum Beispiel für die Beschreibung der *Zeit*. Dass hierfür endliche Körper wie $\mathbb{F}_p = \mathbb{Z}/p\mathbb{Z}$ nicht brauchbar wären, lässt sich wie folgt einsehen.

Proposition 6.4. *Für jeden Körper K existiert eine eindeutige “strukturhaltende Abbildung” $\phi : \mathbb{N} \rightarrow K$ mit $\phi(1) = 1 \in K$. Ist K angeordnet, so ist ϕ injektiv. Insbesondere können angeordnete Körper nicht endlich sein.*

Beweis. Die Idee der “Strukturerhaltung”, wir sagen später *Homomorphismus*, ist, dass wir das gleiche Ergebnis erhalten, egal, ob wir die Rechenoperation vor oder nach dem Anwenden von ϕ ausführen. Hier bedeutet dies also, dass für alle $a, b \in \mathbb{N}$ gilt, dass $\phi(a + b) = \phi(a) + \phi(b)$, wobei das erste Plus in $\mathbb{N}$ und das zweite in K ausgeführt wird. Insbesondere muss gelten: $\phi(n + 1) = \phi(n) + \phi(1)$ für alle n und

daraus wird klar, dass die Abbildung existiert und eindeutig ist.¹⁷ Ist K angeordnet, so folgt aus (6.3) (in K) und den Anordnungsaxiomen (6.2), dass $\phi(n+1) > \phi(n)$ für alle n . Insbesondere ist $\phi(n) \neq \phi(m)$ falls $n \neq m$. ϕ ist also injektiv, und wir können identifizierend $\phi(n) =: n \in K$ schreiben. Da aber $\mathbb{N}$ unendlich ist, kann K nicht endlich sein. Ebenso enthält jeder angeordnete Körper die ganzen Zahlen, und die rationalen Zahlen als angeordneten Unterkörper. $\square$

Ein paar allgemeine Sprachregelungen:

Definition 6.5. Sei $(A, \preccurlyeq)$ eine partiell geordnete Menge. Dann heisst eine Teilmenge $T \subset A$ *nach*¹⁸ *oben* (bzw. *nach unten*) *beschränkt*, falls eine obere (untere) Schranke von T existiert, d.h. ein $B \in A$ so, dass

$$\forall a \in T : a \preccurlyeq B \quad (B \preccurlyeq a) \quad (6.7)$$

Wir sagen *beschränkt* für eine Teilmenge, die sowohl nach unten als auch nach oben beschränkt ist. Falls es eine obere (untere) Schranke gibt, die in T enthalten ist, so ist diese eindeutig und heisst *Maximum* (bzw. *Minimum*) von T , geschrieben $\max(T)$ ($\min(T)$).

Ein Element $m \in T$ heisst *maximal* (*minimal*), falls kein grösseres (kleineres) Element von T existiert, d.h.

$$\forall a \in T : m \preccurlyeq a \Rightarrow a = m \quad (a \preccurlyeq m \Rightarrow a = m) \quad (6.8)$$

Der Unterschied zwischen (6.7) und (6.8) kommt daher, dass wir in einer partiell geordneten Menge nicht unbedingt alle Elemente miteinander vergleichen können.

Übungsaufgabe. Man beweise die Eindeutigkeit des Maximums/Minimums und zeige, dass falls $\preccurlyeq$ total ist, jede endliche nicht-leere Teilmenge stets ein Maximum und ein Minimum besitzt. Zum Unterschied davon sind maximale/minimale Elemente nicht unbedingt eindeutig, existieren für endliche Mengen aber immer.

Lemma 6.6. Ist K ein angeordneter Körper, so ist $T \subset K$ genau dann nach oben durch $B \in K$ beschränkt, falls $-T$ nach unten durch $-B$ beschränkt ist.

Beweis. Aus $a \leq B$ für alle $a \in T$ folgt mittels Lemma 6.2 (iii) $-a \geq -B$ und umgekehrt. $\square$

Beispielsweise sollte die Gesamtenergie eines physikalischen Systems im Sinne der Anordnung nach unten beschränkt sein. Sonst droht (bei Kopplung an externe Freiheitsgrade) Instabilität!

¹⁷Die hierfür nötige Bedingung $\phi(1) = 1$ würde auch aus der Verträglichkeit mit der Multiplikation folgen, d.h. aus $\phi(1) = \phi(1^2) = \phi(1)\phi(1) \Rightarrow \phi(1) = 1$.

¹⁸Auf die Gefahr einer lebenslangen Verwirrung hin sei bemerkt, dass in diesem Zusammenhang “von oben” gleichbedeutend mit “nach oben” ist.

§6. DIE REELLEN ZAHLEN

Definition 6.7. Ein angeordneter Körper K heisst *archimedisch* falls $\mathbb{N} \subset K$ nicht nach oben beschränkt ist. Das heisst $\nexists B \in K : \forall n \in \mathbb{N} : n \leq B$ oder äquivalent:

$$\forall a \in K : \exists N \in \mathbb{N} \text{ s.d. } N > a \quad (6.9)$$

(Hierbei benutzen wir, dass jeder angeordneter Körper die ganzen Zahlen als Teilmenge enthält, wie in Prop. 6.4 festgehalten.)

Bemerkung. · Der Körper der rationalen Zahlen $\mathbb{Q}$ ist archimedisch. Ist $a = \frac{p}{q}$, wobei ohne Einschränkung $p > 0$ und $q \in \mathbb{Z} \setminus \{0\}$, so gilt mit $N = p + 1$: $N > a$.

· Die Konstruktion von angeordneten nicht-archimedischen Körpern leuchtet keinem Anfänger sofort ein. Deshalb ist die Wichtigkeit dieses archimedischen Axioms von allen hier gegebenen Definitionen am schwierigsten einzusehen. Der Ausschluss von Anordnungsunendlichkeiten, welche grösser sind als die natürlichen Zahlen, ist aber wichtig für die Richtigkeit vieler Existenz- und Eindeutigkeitsaussagen der Analysis. Beachte auch, dass der Vergleich von Mengen durch ihre Anordnungseigenschaften zu unterscheiden ist von dem Vergleich über die Mächtigkeit (s. Def. 4.10). Hierzu mehr in der HöMa 2 und 3.

· Aus physikalischer Sicht ist die Unbeschränktheit der natürlichen Zahlen recht natürlich. Jeder einzelne Physiker kommt etwa mit dem Zählen seines Herzschlags bis zum Ende seines Lebens nicht zum Schluss. Für die Menschheit insgesamt besteht Hoffnung, so lange die Erde um die Sonne kreist. Diese Einsicht steht möglicherweise, aber nicht unbedingt, im Konflikt mit der Grundidee der modernen Physik, nach der letztlich alle physikalischen Grössen absoluten Schranken unterliegen (Beispiel: Lichtgeschwindigkeit, Plancksches Wirkungsquantum).

Sowohl in der Analysis als auch in der Physik erfolgen Abschätzungen kleiner Grössen allerdings nicht über “untere Schranken” im Sinne der Anordnung (Def. 6.5), sondern vielmehr über den *Absolutbetrag*, der im Allgemeinen wie folgt definiert ist.

Definition 6.8. Sei K ein angeordneter Körper. Wir setzen für alle $a \in K$:

$$|a| := \max\{a, -a\} = \begin{cases} a & \text{falls } a \geq 0 \\ -a & \text{falls } a < 0 \end{cases} \quad (6.10)$$

Eine Teilmenge $T \subset K$ heisst *beschränkt* (im Sinne des Absolutbetrags), falls ein $B \in K$ existiert so, dass $|a| \leq B$ für alle $a \in T$.

Es gelten die Regeln:

Lemma 6.9.

$$\begin{aligned} \text{Multiplikativität: } & \forall a, b \in K : |a \cdot b| = |a| \cdot |b|, \\ \text{Dreiecksungleichung: } & \forall a, b \in K : |a + b| \leq |a| + |b|, \\ \text{“Umgekehrte” Dreiecksungl.: } & \forall a, b \in K : |a - b| \geq ||a| - |b||, \\ \text{ausserdem: } & |a| = 0 \Leftrightarrow a = 0 \end{aligned} \quad (6.11)$$

Beweis. Die erste Regel verifiziert man leicht anhand einer Fallunterscheidung, die letzte Eigenschaft folgt direkt aus der Definition. Die Dreiecksungleichung folgt wegen $|a + b| = \max\{a + b, -(a + b)\}$ und $a \leq |a|$, $b \leq |b|$ aus der Tatsache, dass eine gemeinsame obere Abschätzung zweier Zahlen auch eine Abschätzung ihres Maximums liefert:

$$\left. \begin{array}{l} a + b \leq |a| + |b| \\ -(a + b) \leq |a| + |b| \end{array} \right\} \Rightarrow \max\{a + b, -(a + b)\} \leq |a| + |b| \quad (6.12)$$

Zum Beweis der umgekehrten Dreiecksungleichung bemerken wir zunächst, dass $|a| = |a - b + b| \leq |a - b| + |b|$ wegen der eben bewiesenen “gewöhnlichen” Dreiecksungleichung. Es gilt also $|a| - |b| \leq |a - b|$. Ebenso gilt $|b| - |a| \leq |a - b|$, zusammen also wieder $||a| - |b|| = \max\{|a| - |b|, |b| - |a|\} \leq |a - b|$. $\square$

Und man überlegt sich leicht:

Lemma 6.10. *Eine Teilmenge $T \subset K$ eines angeordneten Körpers ist genau dann im Sinne des Absolutbetrags beschränkt, wenn sie im Sinne der Anordnung Def. 6.5 (nach oben und unten) beschränkt ist. Die Kurzbezeichnung “beschränkt” ist also in jedem Fall eindeutig.*

$\square$

Die leichte Umformulierung des archimedischen Axioms 6.7 zu der Aussage: “Für alle $B, \epsilon \in K$, $\epsilon > 0$ existiert ein $N \in \mathbb{N}$ s.d. $N\epsilon > |B|$.” lässt sich ebenfalls als Ausdruck der Vergleichbarkeit endlicher (d.h. weder null noch unendlich) physikalischer Größen interpretieren. (Dies war die Motivation von Archimedes.)

Proposition 6.11. *Sei K ein archimedisch angeordneter Körper, und $q, Q \in K$.*

- (i) *Ist $Q > 1$ so gibt es zu jedem $B \in K$ ein $n \in \mathbb{N}$ so, dass $Q^n > B$.*
- (ii) *Ist $0 < q < 1$ so gibt es zu jedem $\epsilon > 0$ ein $n \in \mathbb{N}$ so, dass $q^n < \epsilon$.*

Beweis. (i) Wir können in eindeutiger Weise schreiben $Q = 1 + x$ mit $x > 0$. Die Bernoullische Ungleichung Lemma 6.3 liefert $Q^n \geq 1 + nx$. Weil K archimedisch ist, existiert ein $n \in \mathbb{N}$ mit $nx > B$. Damit gilt $Q^n > B$.

(ii) Folgt aus (i) durch $Q = q^{-1}$ und $B = \epsilon^{-1}$. $\square$

Vollständigkeit

Bevor wir nun zur entscheidenden Charakterisierung der reellen Zahlen kommen, gehen wir noch einmal zurück zur Erweiterung der natürlichen Zahlen zu den rationalen Zahlen durch “Hinzufügen” der Lösungen von (6.1). Es ist leicht einzusehen, dass die Lösbarkeit dieser Gleichungen äquivalent ist zur Aussage:

$$\forall (a_1, a_0) \in K^\times \times K : \exists ! x : a_1 \cdot x + a_0 = 0 \quad (6.13)$$

Derartige Aufgaben sind bekannt als *lineare Gleichungen*. Im Kapitel 3 geht es unter Anderem um die Erweiterung solcher Gleichungen auf mehrere Variablen und ihre allgemeine Lösungstheorie. Andererseits haben, wie hinreichend bekannt, gewisse natürlich auftretende nicht-lineare Aufgaben keine Lösung in $\mathbb{Q}$.

§6. DIE REELLEN ZAHLEN

Theorem 6.12. *Angenommen, $n \in \mathbb{N} \setminus \{m^2 \mid m \in \mathbb{N}\}$. Dann existiert kein $x \in \mathbb{Q}$ so, dass*

$$x^2 = n \quad (6.14)$$

Beweis. Wir erinnern daran, dass sich jede rationale Zahl in eindeutiger Weise schreiben lässt als $x = \frac{p}{q}$ mit $(p, q) \in \mathbb{Z} \times \mathbb{N}$ und teilerfremd. Dann sind auch p^2 und q^2 teilerfremd und $\frac{p^2}{q^2}$ die Darstellung von x^2 . Falls $\frac{p^2}{q^2} = n \in \mathbb{N}$ so ist $q = 1$ und $n = p^2 \in \{m^2 \mid m \in \mathbb{N}\}$. $\square$

Eine Möglichkeit zur Abhilfe ist die weitere Erweiterung der rationalen Zahlen um die fehlenden Lösungen von Gleichungen wie (6.14) und deren Verwandten höheren Grades (s. Gl. (7.3)). Diese Option ist vom mathematischen Standpunkt her ökonomisch und wird in der Algebra gewählt. In dieser Vorlesung folgen wir hingegen der Mutmassung, dass es für die Physik ja ausreichen könnte, Gleichungen wie (6.14) *näherungsweise* zu lösen (auch wenn die physikalische Bedeutung solcher Gleichungen vielleicht zunächst gar nicht klar ist). Die Idee, dass man solche Näherungen “*im Prinzip* beliebig verbessern können muss”, führt ausgehend von den rationalen direkt zu den reellen Zahlen $\mathbb{R}$. Zur Illustration machen wir an dieser Stelle mit der wohlbekannten Aufgabe des Ziehens von Quadratwurzeln eine Klammer auf, die wir erst in Gl. (14.28) mit dem Beweis der Konvergenz des Näherungsverfahrens wieder schliessen werden.

Es sei $x_0 \in \mathbb{Q}$ eine “näherungsweise” Lösung der Gleichung $x^2 = 2$, das heisst $\delta_0 := x_0^2 - 2 \in \mathbb{Q}$ sei “klein” in einem geeigneten Sinne. Der Einfachheit halber nehmen wir an, dass $x_0 > 0$ und $x_0^2 > 2$, also auch $\delta_0 > 0$. (Es ist leicht zu sehen, dass solch ein x_0 existiert.) Definieren wir dann

$$x_1 := x_0 - \frac{\delta_0}{2x_0} \quad (6.15)$$

so gilt

- (i) $x_1 > 0$ ($\Leftrightarrow 0 < 2x_0^2 - \delta_0 = x_0^2 + 2 \checkmark$)
- (ii)

$$\delta_1 := x_1^2 - 2 = \left(x_0 - \frac{\delta_0}{2x_0}\right)^2 - 2 = x_0^2 - 2 - \delta_0 + \left(\frac{\delta_0}{2x_0}\right)^2 = \frac{\delta_0}{x_0^2} \cdot \frac{\delta_0}{4} \quad (6.16)$$

Wegen $\delta_0 = x_0^2 - 2 < x_0^2$ ist damit $0 < \delta_1 < \frac{\delta_0}{4}$ und daher x_1 eine (noch) “bessere” approximative Lösung als x_0 . Wegen (i) und $\delta_1 > 0$ lässt sich diese Prozedur wiederholen. Allerdings ist $x_1 \in \mathbb{Q}$ und keine noch so lange Iteration kann eine exakte Lösung liefern, da ja kein $x \in \mathbb{Q}$ existiert mit $x^2 = 2$. Was tun?

Die reellen Zahlen entstehen durch “Füllen aller Lücken”, die durch solche Iterationen in den rationalen Zahlen aufgetan werden. Das weitere Ziel der Analysis ist dann die Identifikation reeller Zahlen, die durch solche Prozesse auseinander hervorgehen, und die Untersuchung von deren Eigenschaften. Die ursprüngliche “Vollständigkeit der reellen Zahlen” lässt sich dabei auf verschiedene äquivalente Weisen erfassen. Manche davon resonieren besonders gut mit der physikalischen Idee der Approximation (mittels Absolutbetrag), und wir werden sie in Kapitel 4 ausführlich besprechen. Am elegantesten zum Ziel führt jedoch die folgende Formulierung und

Konstruktion, welche nur die Anordnung benutzt, ohne archimedisches Axiom, und auch ohne Folgen auskommt.

Wir setzen den Einschub 6.5 etwas fort.

Definition 6.13. Sei $(A, \preccurlyeq)$ eine partiell geordnete Menge, und $T \subset A$. Dann heisst $B \in A$ *kleinste obere Schranke* von T falls gilt

- (i) B ist eine obere Schranke, d.h. $\forall a \in T : a \preccurlyeq B$ und
- (ii) Für jede andere obere Schranke B' von T gilt, dass $B \preccurlyeq B'$.

Falls eine kleinste obere Schranke existiert, so ist diese eindeutig, und heisst das *Supremum* von T geschrieben: $\sup T$. Der Begriff der grössten unteren Schranke ist entsprechend definiert. Sie heisst *Infimum* von T , geschrieben: $\inf T$.

Übungsaufgabe. Man zeige die Eindeutigkeit und ausserdem: Ist $\sup T \in T$, so ist $\sup T = \max T$, dem Maximum von T , bzw. $\inf T \in T \Rightarrow \inf T = \min T$.

Definition 6.14. Ein angeordneter Körper K heisst (*anordnungs-*)*vollständig*, falls jede nicht-leere nach oben beschränkte Menge eine kleinste obere Schranke (in K !) besitzt.

Bemerkungen. · Wegen Lemma 6.6 wäre die Definition über grösste untere Schranken für nach unten beschränkte Mengen äquivalent.

· Die leere Menge ist zwar trivialerweise beschränkt, kann aber unmöglich eine kleinste obere oder grösste untere Schranke haben. (Warum?)

· In diesem Kontext zeigt das obige Beispiel, dass der Körper der rationalen Zahlen *nicht* anordnungs-vollständig ist: Die Menge $\{x \in \mathbb{Q} \mid x^2 < 2\} \subset \mathbb{Q}$ ist zwar nicht leer und nach oben beschränkt. Sie besitzt aber keine kleinste obere Schranke: Jede obere Schranke $x_0 \in \mathbb{Q}$ erfüllt $x_0 > 0$ und $x_0^2 > 2$, dann ist aber $x_1 = x_0 - \delta_0/(2x_0)$ eine weitere, kleinere, obere Schranke. Alternativ kann man zeigen, dass eine etwaige kleinste obere Schranke dieser Menge eine Lösung der Gleichung $x^2 = 2$ wäre, welche bekanntlich nicht rational ist.

Theorem 6.15. (i) *Es existiert ein bis auf Isomorphismus¹⁹ eindeutig bestimmter vollständiger angeordneter Körper, genannt der Körper der reellen Zahlen $\mathbb{R}$.*

(ii) $\mathbb{R}$ ist archimedisches.

(iii) Jeder angeordnete archimedische Körper ist in $\mathbb{R}$ enthalten.

Einige Beweisideen. (i) Die folgende explizite Beschreibung von $\mathbb{R}$, welche auf R. Dedekind zurückgeht, sichert die Existenz von grössten unteren (bzw. kleinsten oberen) Schranken “per Konstruktion”. Die Idee ist die Identifikation einer *reellen* Zahl mit der Menge derjenigen *rationalen* Zahlen, welche strikt grösser oder kleiner sind.

· Ein *Dedekindscher Schnitt* ist eine Teilmenge $A \subset \mathbb{Q}$ mit den Eigenschaften, dass

(α) $A \neq \emptyset$ und $A \neq \mathbb{Q}$

(β) A ist nach oben ordnungs-abgeschlossen, d.h. $\forall a \in A, x \in \mathbb{Q} : a < x \Rightarrow x \in A$ und

¹⁹Der Begriff des Isomorphismus ist hier noch nicht definiert. Aber wir beweisen ja auch das Theorem nicht.

§6. DIE REELLEN ZAHLEN

(γ) A enthält kein kleinstes Element: $\forall a \in A : \exists x \in A : x < a$

Beispiele für solche Schnitte sind für $r \in \mathbb{Q}$: $A_r = \{x \in \mathbb{Q} \mid x > r\}$. Die wesentliche Beobachtung ist aber, dass es Schnitte von $\mathbb{Q}$ gibt, die nicht von der Form A_r für ein $r \in \mathbb{Q}$ sind, so etwa $\{x \in \mathbb{Q} \mid x > 0 \text{ und } x^2 > 2\}$. Als Menge sind die reellen Zahlen gerade die Menge aller Dedekindschen Schnitte der rationalen Zahlen.

- Die Rechenoperationen sind so konstruiert, dass sie für die A_r mit den üblichen Operationen übereinstimmen, d.h. $A_{r_1} + A_{r_2} = A_{r_1+r_2}$, $A_{r_1} \cdot A_{r_2} = A_{r_1 \cdot r_2}$ etc. Die Anordnung ist definiert über $A < B \Leftrightarrow B \subsetneq A$.

- Die genaue Definition und die Verifikation aller Eigenschaften erfordert einige Arbeit, für die wir auf die Übungen verweisen. Das Endergebnis ist tatsächlich ein vollständiger angeordneter Körper: Eine nicht-leere Menge T Dedekindscher Schnitte ist nach unten beschränkt, wenn $\cup_{A \in T} A \neq \mathbb{Q}$, und in diesem Fall ist $\cup_{A \in T} A$ wieder ein Dedekindscher Schnitt und eine grösste untere Schranke von T . Insbesondere ist bemerkenswert, dass wir durch das Schliessen der Lücken in $\mathbb{Q}$ keine weiteren Lücken aufgetan haben. Zum Beweis der Eindeutigkeit von $\mathbb{R}$ (bis auf Isomorphismus) sagen wir nichts.

(ii) Angenommen, es gäbe ein $a \in K$ so, dass $a \geq N \forall N \in \mathbb{N}$. Dann wäre a eine obere Schranke für $\mathbb{N}$, und es gäbe eine kleinste obere Schranke, a_0 . Nun folgte aber aus $a_0 \geq N$ für alle $N \in \mathbb{N}$ auch $a_0 - 1 \geq N$ für alle $N \in \mathbb{N}$. Dann aber wäre $a_0 - 1 < a_0$ ebenfalls eine obere Schranke, kleiner als a_0 . $\nexists$

(iii) Siehe Ethan Bloch, Real Numbers and Real Analysis für eine sehr benutzerfreundliche und ausführliche Behandlung dieses Zugangs zu den reellen Zahlen. $\square$

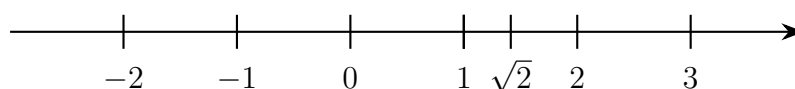
Für das praktische Rechnen sind die Dedekindschen Schnitte natürlich zu unhandlich. Im Folgenden setzen wir die Existenz der reellen Zahlen mit den oben genannten Eigenschaften *voraus*, und formulieren die Vollständigkeit in Kapitel 4 auf so viele verschiedene Weisen um, wie es für unsere Zwecke nützlich ist.

- Es ist letztlich dem Genie Isaac Newtons die Idee zu verdanken, dass die idealisierte Einführung der reellen Zahlen es nicht nur erlaubt, (gewisse, aber bei weitem (noch) nicht alle, siehe nächster §) algebraische Gleichungen zu lösen, sondern auch, Differentialgleichungen zu formulieren, deren Lösungen (“Integrale”) die Bewegungen von Massenpunkten, die Ausbreitung von Wellen und klassischen Feldern, und dergleichen mehr, in einer geeigneten Näherung hervorragend beschreiben.

- Ein Ziel dieser Vorlesungsreihe ist der theoretische Nachweis der Existenz solcher Lösungen in diesem Rahmen sowie die Entwicklung geeigneter Lösungsmethoden.

Visualisierung

Es ist zweckmässig, sich den Bereich der reellen Zahlen mit seiner Anordnung bildlich darzustellen wie aus der Grundschule bekannt. Dabei werden insbesondere Addition, Anordnung und in gewissem Sinne die Vollständigkeit erfasst, die Multiplikation nur durch Anwendung eines Rechenschiebers.



Zwei Ergänzungen

Besonders interessant und als Definitionsbereiche von Funktionen wichtig sind diejenigen Teilmengen von $\mathbb{R}$, die alle zwischen zwei Grenzen liegenden Zahlen enthalten.

Definition 6.16. Ein *Intervall* ist eine “ordnungs-konvexe” Teilmenge I von $\mathbb{R}$. Das heisst, für alle $x, y \in I, \forall z \in \mathbb{R} : x < z < y \Rightarrow z \in I$. Wir unterscheiden für $a \leq b$

das abgeschlossene Intervall: $[a, b] := \{x \mid a \leq x \leq b\}$

die “halb-offenen” Intervalle: $\begin{cases} [a, b) := \{x \mid a \leq x < b\} \\ (a, b] := \{x \mid a < x \leq b\} \end{cases}$

das offene Intervall: $(a, b) := \{x \mid a < x < b\}$

Beachte: Die *unbeschränkten Intervalle* $[a, \infty)$, $(-\infty, a]$ gelten als abgeschlossen, (a, ∞) , $(-\infty, a)$ als offen. Die “Grenze” ∞ ($\notin \mathbb{R}!$) ist also sowohl offen als auch abgeschlossen.²⁰ Beschränkte und abgeschlossene Intervalle heissen *kompakt*.

Proposition 6.17. (i) Zu jeder reellen Zahl $a > 0$ und jeder natürlichen Zahl $n \in \mathbb{N}$ gibt es genau eine reelle Zahl $x > 0$ mit $x^n = a$. Wir schreiben auch $x = a^{1/n}$.

(ii) Aus $b \geq a > 0$ folgt $b^{1/n} \geq a^{1/n}$.

(iii) Es gilt $|x| = (x^2)^{1/2} \forall x > 0$.

Beweis. Lassen wir aus. Man beachte die folgende Variante: Sind $a > 0$ und $b > 0$, so folgt aus $b^n > a^n$, dass auch $b > a$. $\square$

§ 7 Die komplexen Zahlen

Definition 7.1. Die Menge der geordneten Paare $(x, y) \in \mathbb{R} \times \mathbb{R}$ reeller Zahlen, versehen mit den Operationen:

Addition: $(x_1, y_1) + (x_2, y_2) := (x_1 + x_2, y_1 + y_2)$

Multiplikation: $(x_1, y_1) \cdot (x_2, y_2) := (x_1x_2 - y_1y_2, x_1y_2 + x_2y_1)$

und der Auszeichnung $0_{\mathbb{C}} = (0, 0)$ und $1_{\mathbb{C}} = (1, 0)$, ist ein Körper, genannt der Körper der komplexen Zahlen, $\mathbb{C}$.²¹

Bemerkungen. · Von den Körperaxiomen ist nur die Existenz des multiplikativen Inversen etwas nicht-trivial: Für alle $(x, y) \neq (0, 0)$ ist $x^2 + y^2 > 0$, daher

$$(x, y) \cdot \left(\frac{x}{x^2 + y^2}, \frac{-y}{x^2 + y^2} \right) = (1, 0) = 1_{\mathbb{C}} \quad (7.1)$$

²⁰Man beachte auch, dass dem Wort “abgeschlossen” hier ein recht anderer Sinn gegeben wird als in (β). Die ganze Menge $\mathbb{R} = (-\infty, \infty)$ wie auch die leere Menge sind sowohl offen als auch abgeschlossen. Die beschränkten nicht-leeren halb-offenen Intervalle sind weder offen noch abgeschlossen. Ein-Punkt Mengen $\{a\} = [a, a]$ gelten mit dieser Definition auch als Intervalle.

²¹Hier ist ein Beispiel, in dem, anders als unter Def. 2.4 erklärt, beim kartesischen Produkt $\mathbb{C} \times \mathbb{C} = \{(x_1, y_1), (x_2, y_2)\}$ die inneren Klammern nicht aufgelöst werden.

§ 7. DIE KOMPLEXEN ZAHLEN

Gegebenenfalls überprüfe man auch alle anderen Körperaxiome als Übung.

- Die Begründung für die Einführung der komplexen Zahlen liegt in der Möglichkeit der Lösung algebraischer Gleichungen, die in $\mathbb{Q}$ aus *Anordnungsgründen* nicht lösbar sind, und zwar nicht einmal näherungsweise (sodass das Problem auch nicht durch Übergang zu $\mathbb{R}$ gelöst wird). So gilt ja z.B. in jedem angeordneten Körper K , dass $x^2 \geq 0 \forall x \in K$ (s. Lemma 6.2), und daraus folgt $x^2 + 1 \geq 1 > 0$, mithin hat die Gleichung $z^2 + 1 = 0$ keine Lösung in K .

- In $\mathbb{C}$ hingegen gilt

$$(0, 1)^2 = (-1, 0) = -1_{\mathbb{C}} \quad (7.2)$$

und wie wir in Kapitel ?? sehen werden, hat sogar jede nicht-konstante monisch polynomiale Gleichung mit Koeffizienten in $\mathbb{C}$, also

$$z^n + a_{n-1}z^{n-1} + \dots + a_1z + a_0 = 0 \quad a_i \in \mathbb{C}, n \in \mathbb{N} \quad (7.3)$$

eine Lösung $z \in \mathbb{C}$ (Fundamentalsatz der Algebra). (Mit Multiplizität gezählt gibt es dann genau n Lösungen.) Diese *algebraische Abgeschlossenheit von $\mathbb{C}$* ist auch für dessen zentrale Bedeutung in der Quantenmechanik verantwortlich.

- Hingegen ist es als Folgerung aus den obigen Betrachtungen unmöglich, $\mathbb{C}$ in einer Weise anzuordnen, die mit den Körperoperationen verträglich ist. Schreibt man im komplexen Kontext eine Ungleichung wie $B > 0$ so meint dies “ $B \in \mathbb{R} \subset \mathbb{C}$ und positiv im Sinne der Anordnung auf $\mathbb{R}$.” Hierbei ist $\mathbb{R}$ eingebettet in $\mathbb{C}$ via

$$\mathbb{R} \ni x \mapsto (x, 0) \in \mathbb{C} \quad (7.4)$$

d.h., wir identifizieren $x \in \mathbb{R}$ mit $(x, 0) \in \mathbb{C}$.

- Definieren wir dann $i := (0, 1) \in \mathbb{C}$, so können wir jede komplexe Zahl in eindeutiger Weise schreiben als

$$z = (x, y) = (x, 0) + (0, y) = (x, 0) + (0, 1) \cdot (y, 0) = x + iy, \quad x, y \in \mathbb{R} \subset \mathbb{C} \quad (7.5)$$

und “wie gewöhnlich” rechnen. Man sollte sich allerdings möglichst daran gewöhnen, komplexe Zahlen als eine Einheit aufzufassen.

- “Statt” der Anordnung haben wir auf $\mathbb{C}$ die Operation der komplexen Konjugation:

$$x + iy = z \mapsto \bar{z} = x - iy \quad (7.6)$$

Sie erfüllt Rechenregeln wie $\overline{z_1 + z_2} = \bar{z}_1 + \bar{z}_2$, $\overline{z_1 z_2} = \bar{z}_1 \bar{z}_2$, $\bar{\bar{z}} = z$. Real- und Imaginärteil sind

$$\operatorname{Re}(z) := \frac{z + \bar{z}}{2} = x, \quad \operatorname{Im}(z) := \frac{z - \bar{z}}{2i} = y \quad (7.7)$$

und es gilt $\bar{z}z = x^2 + y^2 > 0$ für alle $z \neq 0$. Wir definieren den *Absolutbetrag*

$$|z| := \begin{cases} (\bar{z}z)^{1/2} & z \neq 0 \\ 0 & z = 0 \end{cases} \quad (7.8)$$

welcher wegen Prop. 6.17 auf $\mathbb{R}$ mit (6.10) übereinstimmt. Es gelten $|z| = 0 \Leftrightarrow z = 0$, Abschätzungen der Art

$$|\operatorname{Re}(z)| \leq |z| \quad |\operatorname{Im}(z)| \leq |z|, \quad (7.9)$$

sowie insbesondere die Dreiecksungleichung

$$|z + w| \leq |z| + |w| \quad (7.10)$$

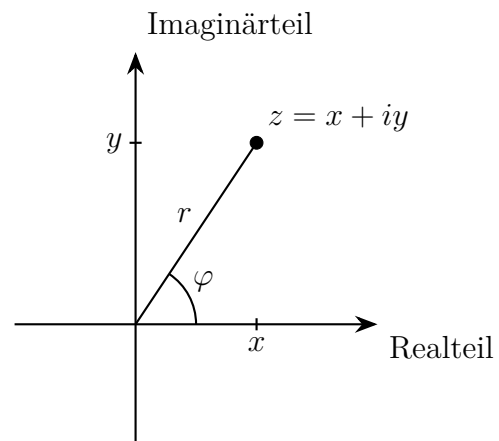
(vgl. Lemma 6.9): Aus den Rechenregeln folgt

$$|z + w|^2 = |z|^2 + |w|^2 + 2 \operatorname{Re}(\bar{z}w) \leq |z|^2 + |w|^2 + 2|z| \cdot |w| = (|z| + |w|)^2 \quad (7.11)$$

unter Benutzung von (7.9). Da für $a \geq 0$, $b \geq 0$ aus $a^2 \geq b^2$ folgt, dass auch $a \geq b$ (vgl. 6.17), impliziert dies (7.10).

Visualisierung

Die Eigenschaften dieser Definitionen erlauben es, sich $\mathbb{C}$ (als Ergänzung zu unserer algebraischen Absicht eben doch) als die Euklidische Ebene vorzustellen, in der die Addition von komplexen Zahlen genau der Addition von Vektoren entspricht. In diesem Bild entspricht komplexe Konjugation der Spiegelung an der reellen Achse, und der Absolutbetrag ist der Abstand vom Ursprung.



• Schreiben wir dann für $z \neq 0$

$$z = r \left(\underbrace{\frac{x}{r}}_{=:c} + i \underbrace{\frac{y}{r}}_{=:s} \right) \quad (7.12)$$

und nehmen die aus der Elementargeometrie bekannte Tatsache vorweg, dass es für zwei reelle Zahlen c und s mit $c^2 + s^2 = 1$ genau einen Winkel $\varphi \in [0, 2\pi)$ gibt mit $c = \cos \varphi$ und $s = \sin \varphi$, so sehen wir, dass

$$z = r(\cos \varphi + i \sin \varphi) \quad (7.13)$$

Mit Hilfe der Additionstheoreme für $\sin$ und $\cos$ folgt dann leicht, dass Multiplikation mit z geometrisch interpretiert werden kann als Drehung um den Winkel φ , gefolgt von Streckung um den Faktor r : Ist $w = s(\cos \chi + i \sin \chi)$ so gilt

$$z \cdot w = r \cdot s (\cos(\varphi + \chi) + i \sin(\varphi + \chi)) \quad (7.14)$$

Achtung: (7.14) ist im Allgemeinen *nicht* die Darstellung von $z \cdot w$ in der Form (7.13), denn $\varphi + \chi$ kann ausserhalb des Intervalls $[0, 2\pi)$ liegen. Streng genommen machen diese Betrachtungen auch erst nach der Definition der Winkelfunktionen und dem Nachweis ihrer Periodizität Sinn.

• Daraus folgt z.B. mittels vollständiger Induktion dann auch die *Formel von de-Moivre*

$$z^n = r^n (\cos(n\varphi) + i \sin(n\varphi)) \quad (7.15)$$

sowie die $\mathbb{C}$ -Version von Prop. 6.17, einem Spezialfall von (7.3).

§ 8. PHYSIKALISCHE RÄUME

Proposition 7.2. Für jede komplexe Zahl $a \neq 0$ und jede natürliche Zahl n existieren genau n Lösungen der Gleichung

$$z^n = a$$

Beweis. Für $a = r(\cos \varphi + i \sin \varphi)$ lösen

$$z_k = r^{1/n} \left(\cos \left(\frac{\varphi}{n} + \frac{2\pi k}{n} \right) + i \sin \left(\frac{\varphi}{n} + \frac{2\pi k}{n} \right) \right) \quad (7.16)$$

für alle $k = 0, 1, \dots, n-1$ die gegebene Gleichung. Dass dies alle Lösungen sind, folgt aus der aus den Übungen bekannten Tatsache, dass ein Polynom vom Grad n höchstens n verschiedene Wurzeln besitzt. $\square$

· Wir bemerken zum weiteren Vergleich mit $\mathbb{R}$ noch, dass wir die Vollständigkeit von $\mathbb{C}$ (die wir intuitiv aus der geometrischen Darstellung als $\mathbb{R} \times \mathbb{R}$ natürlich schon jetzt erfassen) wegen der fehlenden Anordnung *nicht* über die Supremumseigenschaft ausdrücken können, sondern erst über eine der im Kapitel 4 gegebenen Formulierungen. Für den Beweis des Fundamentalsatzes der Algebra spielt dann die Vollständigkeit eine wesentliche Rolle.

· Obwohl also in der komplexen Welt Aussagen wie “ $z > w$ ” im Allgemeinen keinen Sinn machen (vielmehr impliziert man für gewöhnlich mit so einer Schreibweise, dass $z, w \in \mathbb{R} \subset \mathbb{C}$), so können wir dennoch vereinbaren, dass

Definition 7.3. Eine Teilmenge $T \subset \mathbb{C}$ heisst beschränkt, falls eine Schranke $B \in \mathbb{R}$ existiert mit

$$|z| < B \quad \forall z \in T$$

Damit sind das Analogon der Intervalle als “elementare” Definitionsbereiche von Funktionen im Komplexen die *Kreisscheiben* vom Radius $R > 0$ um den Punkt $z \in \mathbb{C}$:

$$D_R(z) := \{w \in \mathbb{C} \mid |z - w| < R\} \quad (7.17)$$

sind “offen”,

$$\overline{D_R(z)} := \{w \in \mathbb{C} \mid |z - w| \leq R\} \quad (7.18)$$

“abgeschlossen”.

§ 8 Physikalische Räume

Dieser § erfüllt einen doppelten Zweck. Er dient einerseits zur Einstimmung auf die mathematischen Strukturen (als Versuch zu deren physikalischer Begründung), mit denen wir uns im restlichen Semester primär beschäftigen werden: Lineare Algebra, Vollständigkeit und Stetigkeit. Hierzu stellt er andererseits einige Begriffe der klassischen Vektorrechnung zusammen, um damit die Verbindung zu den physikalischen Grundvorlesungen zu erleichtern.

In unserer Auffassung besteht die Naturbeschreibung (durch die Physik wie auch die übrigen Naturwissenschaften) aus dem *Hervorzeigen* gewisser realer Sachverhalte (Objekte und Ereignisse), der *Zusammenfassung* solcher Sachverhalte zu Mengen im Sinne der Vereinbarung 2.1 und dem anschliessenden *Vergleich* mit geeigneten

mathematischen Strukturen, die sich wie angedeutet prinzipiell stets als Teilmengen mit gewissen Eigenschaften definieren lassen.

Ein vertrautes Beispiel, auf das wir bereits mehrfach angespielt haben, ist die Benutzung einer Anordnung im Sinne von Def. 3.3 für die Erfassung der zeitlichen Abfolge von “Ereignissen” (durch gegebene “Beobachter” an einem “festen Ort”²²). In einer quantitativen Beschreibung, die man durch Abzählen und Vergleichen “regelmässig wiederkehrender” Ereignisse gewinnt, ergibt sich daraus die reelle Zahlenreihe als Bild für die physikalische Zeit.²³ Man beachte allerdings:

1. Die Rechenoperationen machen Sinn *nicht* zur Verknüpfung verschiedener Zeitpunkte, sondern nur als *Operationen* auf der Menge der Ereignisse.²⁴
2. Es ist kein Zeitpunkt als ‘0’ ausgezeichnet. (Selbst im kosmologischen Kontext zeichnet der Urknall keinen Punkt auf der reellen Zeitachse aus, sondern macht vielmehr das Bild als solches fundamental unbrauchbar.)
3. Mengen physikalischer Sachverhalte sind stets *endlich*. Deshalb ist auch die Vollständigkeit von $\mathbb{Q}$ zu $\mathbb{R}$ nur eine Idealisierung ohne operationelles Gegenstück in der Natur.

Als weiteres Beispiel soll kurz skizziert werden, wie man beginnen könnte, die für die Beschreibung des physikalischen Raums relevanten Strukturen in ähnlicher Weise aus der Erfahrung heraus operationell zu begründen. Am Anfang der klassischen Mechanik wird üblicherweise ohne Weiteres die Aussage vorausgesetzt:

“Der physikalische Raum lässt sich mathematisch als drei-dimensionaler euklidischer Raum beschreiben.” (8.1)

Tatsächlich ist die angeführte Struktur (des drei-dimensionalen euklidischen Raums) aus mehreren Komponenten aufgebaut, von denen jede für sich bereits eine umfangreiche mathematische Theorie verdient. Aus physikalischer Sicht ist eine solche *Zerlegung* zunächst weder einfach noch wirklich notwendig. Allerdings kann die Reflexion über die einzelnen *Hypothesen* zu neuen physikalischen Einsichten führen (z.B. im Kontext von Einsteins Allgemeiner Relativitätstheorie) und auch wieder in völlig unerwarteter Weise an anderer Stelle nützlich werden (z.B. im Kontext der Quantenmechanik).

Man könnte also starten²⁵ von der rein qualitativen Beobachtung, dass physikalische Objekte oder “Körper” eine räumliche Ausdehnung haben und zueinander in der Beziehung des “Berührens” oder “Aneinanderstossens” stehen können. Sie können sodann dazu verwendet werden, “Raumgebiete auszufüllen”.²⁶ Ist Z eine solche

²²Die Bedeutung dieser angeführten Begriffe lässt sich wieder nur durch Vereinbarung festlegen.

²³Aus diesem Grund werden wir in der Diskussion der Differentialrechnung (zu der wir in diesem Durchlauf nicht im ersten Semester kommen) stets ‘ t ’ als Symbol für eine ein-dimensionale unabhängige Variable benutzen, bei der es auf die Anordnung ankommt.

²⁴Für die Multiplikation folgt dies bereits aus Dimensionsbetrachtungen, bei der Addition aus dem zweiten Punkt. Vielmehr entspricht die Addition einer “Verschiebung” um bestimmte Zeitintervalle, die Multiplikation der Wiederholung solcher Verschiebungen, oder dem Strecken der Zeitachse durch Wechsel der Masseinheiten.

²⁵Das folgende Gedankenexperiment stammt nicht von mir. Wem dies zu weit geht, möge statt dessen sich ins Gedächtnis rufen, wie man operationell den rechten Winkel definiert. Das geht ja wohl nicht ohne Auszeichnung von Punkten, Geraden, und Abständen!

²⁶Je nach Material und Geometrie mit mehr oder weniger grossen Lücken. Eine denkbare Variante ist das Aufteilen von Raumgebieten mit Mauern und Wänden.

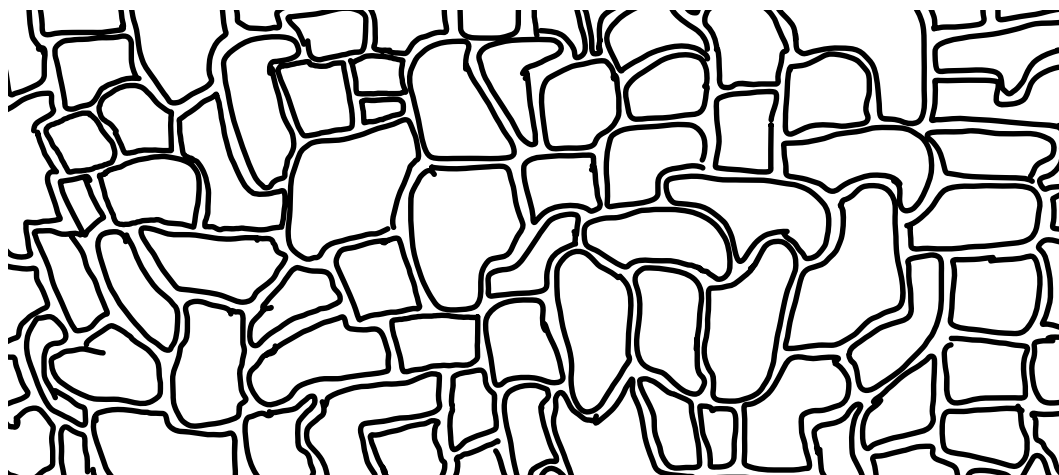


Abbildung 8.1: Bei einer vollständigen Aufteilung einer zwei-dimensionalen Welt kann es Vier-, Fünf- und sogar Siebzehnländerecke geben, muss aber nicht. Dreiländerecke sind unvermeidbar.

Füllung (Zerlegung) eines Raumgebiets G , so können wir an jeder²⁷ Stelle $p \in G$ zählen, wie viele Körper von Z in p aneinanderstossen. Nennen wir diese Zahl $\beta_Z(p)$, so gäbe es stets Stellen mit $\beta_Z(p) = 1$ (im Innern der Körper), Stellen mit $\beta_Z(p) = 2$, 3 und 4, und manchmal auch Stellen mit $\beta_Z(p) = 5$ oder mehr, die man allerdings durch kleine Umbauten eliminieren könnte. Mit anderen Worten, ist $\mathcal{Z} \ni Z$ eine Menge derartiger Zerlegungen, so würden wir finden

$$\min_{Z \in \mathcal{Z}} \left(\max_{p \in G} \beta_Z(p) \right) = 4 \quad (8.2)$$

und hätten damit die Dimension unserer makroskopischen Welt als $4 - 1 = 3$ “entdeckt”. (Wäre unsere Welt etwa zweidimensional, so wäre diese Zahl 3, siehe Bild 8.1.) Weitere Struktur ergäbe sich durch Zählen, Vergleichen und Verkleinern der Körper, natürlich unter Zuhilfenahme moderner Technologie, was wir aber nicht im Einzelnen ausführen. Stattdessen beschreiben wir, “das Pferd von hinten aufzäuhend”, das mathematische Ergebnis unter Vorwegnahme einiger Definitionen der folgenden Kapitel, und zwar direkt für $3 = n \in \mathbb{N}$.

Definition 8.1. Der n -dimensionale *reelle lineare Raum* ist die Menge der (als Spaltenvektoren geschriebenen) n -Tupel

$$\mathbb{R}^n = \left\{ v = \begin{pmatrix} v^1 \\ v^2 \\ \vdots \\ v^n \end{pmatrix} \mid v^i \in \mathbb{R} \text{ für } i = 1, 2, \dots, n \right\} \quad (8.3)$$

²⁷In der Praxis natürlich nur an endlich vielen.

zusammen mit den Operationen

$$\begin{aligned} \text{Vektoraddition: } \mathbb{R}^n \times \mathbb{R}^n \ni (v_1, v_2) &\mapsto v_1 + v_2 = \begin{pmatrix} v_1^1 + v_2^1 \\ \vdots \\ v_1^n + v_2^n \end{pmatrix} \in \mathbb{R}^n \\ \text{Skalarmultiplikation: } \mathbb{R} \times \mathbb{R}^n \ni (\alpha, v) &\mapsto \alpha \cdot v = \begin{pmatrix} \alpha v^1 \\ \vdots \\ \alpha v^n \end{pmatrix} \in \mathbb{R}^n \end{aligned} \quad (8.4)$$

Diese Verknüpfungen genügen den offensichtlichen Regeln (Assoziativität, Kommutativität und Invertierbarkeit der Addition mit Nullvektor $v^i = 0$ für alle i , Assoziativität der Skalarmultiplikation über der Multiplikation von Skalaren, triviale Wirkung von $\alpha = 1$, sowie die zwei möglichen Distributivgesetze), und machen den $\mathbb{R}^n$ zu einem n -dimensionalen Vektorraum über $\mathbb{R}$ im Sinne der Defs. 9.1 und 9.21. In der konkreten Deutung ist der $\mathbb{R}^n$ der “Raum der Richtungen”, in die man von einem Punkt des Raumes startend schauen oder laufen kann. Wohlgermerkt impliziert die gegebene Formulierung nicht, dass die Richtungen, die durch die *Basisvektoren*

$$e_i = \begin{pmatrix} 0 \\ \vdots \\ 1 \\ \vdots \\ 0 \end{pmatrix} \quad i\text{-te Zeile} \quad (8.5)$$

mit $v^i = 1$ und 0 sonst gegeben sind, orthogonal zueinander sind (siehe unten, Def. 8.4). Die nicht-triviale Strukturaussage ist lediglich, dass solche Richtungen sich “untereinander addieren” und “strecken” lassen. Sie ist verträglich mit der gesamten aktuellen experimentellen Datenlage.²⁸ Dass man durch Fortsetzen der Richtungen genau jeden anderen Punkt des Raumes erreicht (eine Aussage, die spätestens auf kosmologischen Skalen revidiert werden muss), beschreibt man folgendermassen.

Definition 8.2. Ein n -dimensionaler *reeller affiner Raum* ist eine Menge X zusammen mit einer Abbildung $\tau : \mathbb{R}^n \times X \rightarrow X$, geschrieben $(v, x) \mapsto \tau_v(x)$ mit den Eigenschaften, dass

- (i) $\forall v_1, v_2 \in \mathbb{R}^n : \tau_{v_1} \circ \tau_{v_2} = \tau_{v_1+v_2}$
- (ii) $\forall x_1, x_2 \in X : \exists! v \in \mathbb{R}^n : \tau_v(x_1) = x_2$

Ist (X nicht leer und) $x_0 \in X$ fest gewählt, so ist gemäss Bedingung (ii) jedem $x \in X$ ein eindeutiges $v \in \mathbb{R}^n$ zugeordnet mit der Eigenschaft, dass $\tau_v(x_0) = x$. Wir schreiben für dieses v $\delta(x_0, x)$.

Lemma 8.3. Für jedes $x_0 \in X$ sei $\delta(x_0, \cdot) : X \rightarrow \mathbb{R}^n$ die Umkehrabbildung der gemäss Def. 8.2 (ii) bijektiven Abbildung $\tau_{\cdot}(x_0) : \mathbb{R}^n \rightarrow X$. Dann gilt

$$\forall x_1, x_2, x_3 \in X : \quad \delta(x_1, x_2) + \delta(x_2, x_3) = \delta(x_1, x_3) \quad (8.6)$$

Ausserdem ist für alle $v \in \mathbb{R}^n$ die Abbildung $\tau_v : X \rightarrow X$ bijektiv und es gilt $\tau_0 = \text{id}_X$.

²⁸lässt sich allerdings tbomk nicht ohne Längen- und Winkelmessung operationell erklären.

§ 8. PHYSIKALISCHE RÄUME

Beweis. Dass $\tau_0 = \text{id}_X$ folgt wegen $\tau_0(\tau_v(x_0)) = \tau_{0+v}(x_0) = \tau_v(x_0)$ aus der Tatsache, dass für beliebiges x_0 $\tau_{\mathbb{R}^n}(x_0) = X$. Dann aber ist $\tau_{-v} \circ \tau_v = \tau_v \circ \tau_{-v} = \tau_0 = \text{id}_X$, d.h. τ_v ist invertierbar. Die Gleichung (8.6) folgt dann aus der Tatsache, dass $\delta(x_1, x_3)$ für alle x_1, x_3 eindeutig dadurch bestimmt ist, dass $\tau_{\delta(x_1, x_3)}(x_1) = x_3$. Denn genau diese Bedingung wird wegen (i) für jedes x_2 auch von $\delta(x_1, x_2) + \delta(x_2, x_3)$ erfüllt:

$$\tau_{\delta(x_1, x_2) + \delta(x_2, x_3)}(x_1) = \tau_{\delta(x_2, x_3)}(\tau_{\delta(x_1, x_2)}(x_1)) = \tau_{\delta(x_2, x_3)}(x_2) = x_3 \quad (8.7)$$

□

In der physikalischen Interpretation ist x_0 der “Standpunkt des Beobachters” als beliebig gewählter “Ursprung des Koordinatensystems”. Der Richtungsvektor $v = \delta(x_0, x)$, der von x_0 zu x führt, ist auch bekannt als “Orts-” oder “Abstandsvektor” von x . Die Aussage, “der physikalische Raum ist ein 3-dimensionaler reeller affiner Raum” bedeutet ähnlich wie bei Bem. 2 auf S. 36 (siehe auch die Warnung zu Äquivalenzklassen auf S. 15), dass man (ohne sonstige Erkenntnisse) keinen Punkt des Raums vor den anderen auszeichnen kann. Die Identifikation des physikalischen Raums mit dem $\mathbb{R}^3$ hängt aber nicht nur von der Wahl eines Koordinatenursprungs ab, sondern auch von der oben angesprochenen Auszeichnung der Basisrichtungen e_i . Hierfür benötigt man im mathematischen Bild ein weiteres Strukturelement.

Definition 8.4. Das *euklidische innere “Standard-”Produkt* auf dem $\mathbb{R}^n$ ist die Abbildung

$$\begin{aligned} \langle \cdot, \cdot \rangle : \mathbb{R}^n \times \mathbb{R}^n &\rightarrow \mathbb{R} \\ (v_1, v_2) &\mapsto \langle v_1, v_2 \rangle := \sum_{i=1}^n v_1^i v_2^i \end{aligned} \quad (8.8)$$

Lemma 8.5. Es gilt für alle $v, v_1, v_2, v_3 \in \mathbb{R}^n$, $\alpha \in \mathbb{R}$

$$\begin{aligned} \text{Symmetrie:} \quad &\langle v_1, v_2 \rangle = \langle v_2, v_1 \rangle \\ \text{(Bi-)Linearität:} \quad &\langle v_1 + \alpha v_2, v_3 \rangle = \langle v_1, v_3 \rangle + \alpha \langle v_2, v_3 \rangle \\ \text{Positiv Definitheit:} \quad &\langle v, v \rangle > 0 \quad \text{falls } v \neq 0 \end{aligned} \quad (8.9)$$

Beweis. Die ersten beiden Eigenschaften folgen aus den Rechenregeln, die dritte aus der Positivität von Quadratzahlen, s. Gl. 6.3. □

Die Verbindung zur Winkel- und Längenmessung entsteht wie folgt.

Definition 8.6. (i) Die (euklidische) *Länge* eines Vektors $v \in \mathbb{R}^n$ ist die nicht-negative reelle Zahl

$$\|v\| := \langle v, v \rangle^{1/2} = \left(\sum_{i=1}^n (v^i)^2 \right)^{1/2} \quad (8.10)$$

“ n -dimensionaler Satz des Pythagoras”. Beachte im Fall $n = 1$ auch Prop. 6.17 (iii). (ii) Der Winkel zwischen zwei von null verschiedenen Vektoren $v_1, v_2 \in \mathbb{R}^n \setminus \{0\}$ ist die durch die Gleichung

$$\cos(\sphericalangle(v_1, v_2)) = \frac{\langle v_1, v_2 \rangle}{\|v_1\| \cdot \|v_2\|} \quad (8.11)$$

eindeutig definierte Zahl $\sphericalangle(v_1, v_2) \in [0, \pi)$

Dass sich hier (ähnlich wie bei der Besprechung der komplexen Zahlenebene auf S. 34) Gl. 8.11 nach $\angle(v_1, v_2) \in [0, \pi)$ auflösen lässt, hängt davon ab, dass die rechte Seite im Intervall $[-1, 1]$ liegt. Dies ist gleichbedeutend mit der folgenden, immer wieder benutzten Abschätzung.

Proposition 8.7. *Das euklidische innere Produkt erfüllt die “Cauchy-Schwarz-Ungleichung”:*

$$\forall v_1, v_2 \in \mathbb{R}^n : \quad |\langle v_1, v_2 \rangle| \leq \|v_1\| \cdot \|v_2\| \quad (8.12)$$

Beweis. Für $v_2 = 0$ ($\Leftrightarrow \|v_2\| = 0$) ist die Aussage klar. Für $v_2 \neq 0$, beliebiges v_1 , sei

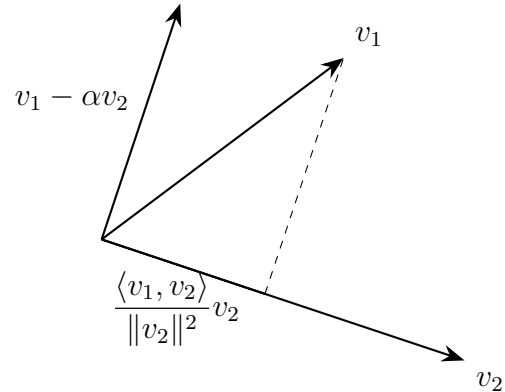
$$\alpha := \frac{\langle v_1, v_2 \rangle}{\|v_2\|^2} \quad (8.13)$$

Dann ist $\langle v_1 - \alpha v_2, v_2 \rangle = 0$ und es gilt Deutung als “orthogonale Projektion”:

$$\begin{aligned} 0 &\leq \|v_1 - \alpha v_2\|^2 \\ &= \|v_1\|^2 + \alpha^2 \|v_2\|^2 - 2\alpha \langle v_1, v_2 \rangle \\ &= \|v_1\|^2 + \frac{\langle v_1, v_2 \rangle^2}{\|v_2\|^2} - 2 \frac{\langle v_1, v_2 \rangle^2}{\|v_2\|^2} \\ &= \|v_1\|^2 - \frac{\langle v_1, v_2 \rangle^2}{\|v_2\|^2} \end{aligned}$$

und damit

$$\langle v_1, v_2 \rangle^2 \leq \|v_1\|^2 \cdot \|v_2\|^2 \quad (8.14)$$



Durch Wurzelziehen folgt mit Prop. 6.17 die Behauptung. $\square$

Die C.-S.U. geht insbesondere ein in den Nachweis der Dreiecksungleichung für die euklidische Länge.

Proposition 8.8. *Es gilt für alle $v, v_1, v_2 \in \mathbb{R}^n$, $\alpha \in \mathbb{R}$:*

- (i) *Homogenität:* $\|\alpha v\| = |\alpha| \cdot \|v\|$
- (ii) *Positivität:* $\|v\| > 0 \quad \forall v \neq 0$
- (iii) *Dreiecksungleichung:* $\|v_1 + v_2\| \leq \|v_1\| + \|v_2\|$

Beweis. Die beiden ersten Eigenschaften sind wieder offensichtlich. Die dritte folgt durch Wurzelziehen aus

$$\begin{aligned} \|v_1 + v_2\|^2 &= \|v_1\|^2 + \|v_2\|^2 + 2\langle v_1, v_2 \rangle \leq \\ &(\text{Benutze hier Prop. 8.7}) \leq \|v_1\|^2 + \|v_2\|^2 + 2\|v_1\| \cdot \|v_2\| = (\|v_1\| + \|v_2\|)^2 \end{aligned} \quad (8.15)$$

$\square$

Die Folgerungen von Prop. 8.8 sind die definierenden Eigenschaften einer *Norm* (s. Def. 15.2). Der Begriff spielt eine zentrale Rolle für die Feststellung der *Vollständigkeit* des $\mathbb{R}^n$, die man ab $n = 2$ nicht mehr durch die Anordnung erklären kann.

§ 8. PHYSIKALISCHE RÄUME

Die Verbindung zur *physikalischen Abstandsmessung*, die wir im Zusammenhang der Abbildung 8.1 angedeutet hatten, entsteht über die affine Struktur aus Def. 8.2 und Lemma 8.3.

Definition 8.9. Ist X ein affiner euklidischer Raum, so ist der euklidische Abstand zweier Punkte $x_1, x_2 \in X$ definiert als Länge des Abstandsvektors,

$$d_E(x_1, x_2) = \|\delta(x_1, x_2)\| \quad (8.16)$$

Nach Wahl eines rechtwinkligen Koordinatensystems können wir mit Hilfe von $\tau.(x_0)$ bzw. $\delta(x_0, \cdot)$, X mit $\mathbb{R}^n$ identifizieren und schreiben

$$d_E(x_1, x_2)^2 = \langle x_2 - x_1, x_2 - x_1 \rangle \quad (8.17)$$

Aus den Eigenschaften (ii) und (iii) von Prop. 8.8 folgt, dass der euklidische Abstand eine “Metrik” ist im Sinne der Def. 15.1.

Lemma 8.10. *Es gelten*

$$\begin{aligned} \text{Symmetrie:} \quad d_E(x_1, x_2) &= d_E(x_2, x_1) \quad \forall x_1, x_2 \in X \\ \text{Positivität:} \quad d_E(x_1, x_2) &= 0 \Leftrightarrow x_1 = x_2 \\ \text{Dreiecksungleichung:} \quad d_E(x_1, x_2) &\leq d_E(x_1, x_3) + d_E(x_3, x_2) \quad \forall x_1, x_2, x_3 \in X \end{aligned} \quad (8.18)$$

Zuletzt könnten wir von der Abstandsmessung wieder abstrahieren und die Aufteilung mit “beliebig geformten” und “verschieden grossen” ausgedehnten Körpern, die an verschiedenen Stellen aneinanderstossen²⁹, als Ausdruck der zugrundeliegenden *topologischen Struktur* des physikalischen Raums erkennen. Zur Existenz und Wahl rechtwinkliger, linearer, und krummliniger Koordinatensysteme später mehr.

Vielleicht nicht unerwähnt bleiben sollte noch das häufig benutzte sogenannte “Vektor-” oder “Kreuzprodukt”. Es wird üblicherweise definiert als Verknüpfung

$$\begin{aligned} \mathbb{R}^3 \times \mathbb{R}^3 &\rightarrow \mathbb{R}^3 \\ (v_1, v_2) &\mapsto v_1 \times v_2 := \begin{pmatrix} v_1^2 v_2^3 - v_1^3 v_2^2 \\ v_1^3 v_2^1 - v_1^1 v_2^3 \\ v_1^1 v_2^2 - v_1^2 v_2^1 \end{pmatrix} \end{aligned} \quad (8.19)$$

und kann (falls v_1 und v_2 nicht kollinear, insbesondere nicht 0 sind) geometrisch interpretiert werden als derjenige Vektor,

1. welcher senkrecht auf v_1 und v_2 steht, d.h.

$$\langle v_1, v_1 \times v_2 \rangle = \langle v_2, v_1 \times v_2 \rangle = 0 \quad (8.20)$$

2. dessen Länge

$$\|v_1 \times v_2\| = \sqrt{\|v_1\|^2 \|v_2\|^2 - \langle v_1, v_2 \rangle^2} = \|v_1\| \cdot \|v_2\| \cdot \sin(\angle(v_1, v_2)) \quad (8.21)$$

²⁹Im mathematischen Bild arbeitet man lieber mit überlappenden sog. “offenen” Mengen, s. HöMa 2&3.

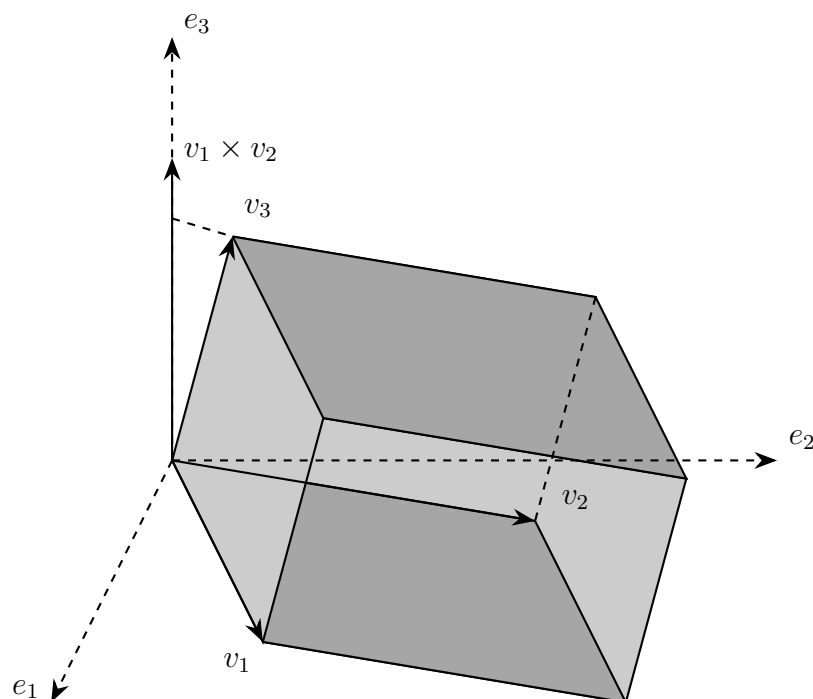


Abbildung 8.2: Das Spatprodukt berechnet das orientierte euklidische Volumen eines Parallelotops.

gleich ist dem *Flächeninhalt* des von v_1 und v_2 aufgespannten Parallelogramms, 3. und der zusammen mit den beiden Eingaben, in der Reihenfolge $(v_1, v_2, v_1 \times v_2)$ ein *positives* (rechtshändig orientiertes) *Koordinatensystem* erzeugt.

Allerdings, und dies ist der wesentliche Grund für diesen Einschub,

1. kann das Kreuzprodukt (jedenfalls so wie in Gl. (8.19)) nur für $n = 3$ definiert werden;
2. ist die Interpretation der Länge als Flächeninhalt aus Dimensionsbetrachtungen verdächtig;
3. ist die “Positivität” bzw. Orientierung ein weiteres Strukturelement, über das wir noch nicht geredet haben.

Etwas besser ist die Verwendung des Kreuzprodukts zur Definition/Berechnung des *orientierten Volumens* eines von drei Vektoren v_1, v_2, v_3 aufgespannten Parallelotops (das Spatprodukt)

$$\begin{aligned} \text{vol}(v_1, v_2, v_3) &= \langle v_1 \times v_2, v_3 \rangle \\ &= v_1^1 v_2^2 v_3^3 + v_1^2 v_2^3 v_3^1 + v_1^3 v_2^1 v_3^2 - v_1^1 v_2^3 v_3^2 - v_1^3 v_2^2 v_3^1 - v_1^2 v_2^1 v_3^3 \end{aligned} \quad (8.22)$$

und dies wird relevant werden als Beispiel für Determinanten in § 10, siehe auch Kapitel ??

KAPITEL 3

VEKTOREN

Wir entwickeln in diesem Kapitel die “abstrakte” Theorie endlich-dimensionaler Vektorräume. Zum Unterschied von § 8 geben wir uns nicht “Tupel reeller Zahlen” vor, welche relativ zu einem Koordinatensystem die Positionen verschiedener physikalischer Objekte spezifizieren, und auf denen wir anschliessend Rechenoperationen definieren. Vielmehr geht es darum, zu verstehen, inwiefern wir *ausgehend von den Rechenoperationen* die Existenz solcher linearer Koordinaten zeigen und die verschiedenen Wahlmöglichkeiten parametrisieren können. Dies führt zu den Begriffen der *Basis* und der *Dimension* eines Vektorraums. Die wesentliche Erkenntnis wird sein, dass Vektorräume nicht unbedingt natürlicherweise mit einer Basis “geboren” werden, man zum konkreten Rechnen (mit Matrizen und linearen Gleichungssystemen) aber stets eine Basis wählen kann und muss, um im Anschluss sicherzustellen, dass die Ergebnisse nicht von der Wahl abhängen. In anderen Worten heisst die Lösung: “Ein Vektorraum (im mathematischen Sinne, also ohne ausgezeichnete Basis) ist ein Vektorraum (im physikalischen Sinne) mit Basis, in dem man die Wahl der Basis vergessen hat oder jedenfalls will.”

Zum Abschluss des Kapitels diskutieren wir weitere Konstruktionen von Vektorräumen, die diese Einsichten noch einmal auf den Punkt bringen.

In Vorbereitung erinnern und erweitern wir leicht unsere Konventionen zum Umgang mit Indexmengen, in zwei Richtungen.

Gemäss Vereinbarungen am Ende von Kapitel 1 (Def. 4.9, siehe auch Gl. (4.7)) ist für eine beliebige Indexmenge I eine “ I -indizierte Familie von Objekten der Sorte M ” einfach eine Abbildung $m_- : I \rightarrow M$. Dieses Konzept erlaubt es (im Unterschied zu den gewöhnlichen Teilmengen) nicht nur, bei der Auswahl manche Elemente von M mehrmals zu berücksichtigen. Vielmehr kann man versuchen, möglicherweise auf I erklärte Strukturen auf die Familie zu “übertragen”. Dies spielt zuvorderst eine Rolle wenn man auf Abzählung aus ist, d.h. wenn I eine (endliche oder unendliche) Teilmenge der natürlichen Zahlen ist. Wir schreiben in solchen Fällen $(m_1, m_2, \dots, m_n)$ (und sagen *endliche Folge* oder *n -Tupel*) oder $(m_k)_{k \in \mathbb{N}} = (m_1, m_2, \dots)$ (und sagen *Folge*), und zwar auch dann, wenn die Familie injektiv ist und es zunächst gar nicht auf die Reihenfolge ankommt, mit anderen Worten wenn wir uns eigentlich primär für die Menge $\{m_1, m_2, \dots\}$ interessieren. Man redet in vielen Fällen wohl manchmal auch gleichbedeutend von *Systemen*. Ein anderes Beispiel ist der Begriff der *Teilfamilie*, formal definiert als Einschränkung der Abbildung m_- auf eine Teilmenge $J \subset I$.

Ist $M = (H, +)$ eine abelsche Halbgruppe (d.h. $+$ ist assoziative und kommutativ)

tiv, möglicherweise ohne neutrales Element und Inverse, s. Def. 5.1), so sind Summen $\sum_{i \in I} a_i$ von endlichen Familien (d.h. $\#I \in \mathbb{N}$) definiert gemäss Gl. (5.1). Ist I eine unendliche Menge, so hat dies im Allgemeinen keinen Sinn, es sei denn, die Operationen mit den allermeisten Mitgliedern ist neutral.

Definition. Ist $(H, +)$ eine abelsche Halbgruppe³⁰ mit Eins 0_H , und X eine beliebige Menge, so ist der (mengentheoretische) *Träger* einer Abbildung $f : X \rightarrow H$ die Teilmenge³¹

$$\text{supp}(f) := \{x \in X \mid f(x) \neq 0_H\}$$

f heisst *von endlichem Träger*, falls $|\text{supp}(f)| < \infty$. Man sagt auch “ $f(x) = 0_H$ für fast alle $x \in X$ ”. Für eine I -indizierte Familie $(h_i)_{i \in I}$ von Elementen von H von endlichem Träger ist die Summe via

$$\sum_{i \in I} h_i := \sum_{i \in \text{supp}(h)} h_i$$

(wohl-)definiert.

Beachte, dass bei einer I -indizierten Familie nicht alle h_i verschieden sein müssen. Insbesondere ist Endlichkeit des Trägers nicht gleichbedeutend mit Endlichkeit des Bildes als Teilmenge von H .

§ 9 Vektorräume

Definition 9.1. Sei K ein Körper. Ein K -Vektorraum oder *Vektorraum über K* ist eine abelsche Gruppe $(V, +_V)$, d.h. es existiert ein $0_V \in V$ so dass für alle $v, v_1, v_2, v_3 \in V$ gilt

- (i) $v_1 +_V v_2 = v_2 +_V v_1$
- (ii) $(v_1 +_V v_2) +_V v_3 = v_1 +_V (v_2 +_V v_3)$
- (iii) $0_V +_V v = v$
- (iv) $\exists -v \in V : v +_V (-v) = 0_V$

mit den üblichen Folgerungen (dass nämlich 0_V und $-v$ eindeutig sind), zusammen mit einer Operation³² $\cdot : K \times V \rightarrow V$, $(\alpha, v) \mapsto \alpha \cdot v$ so, dass für alle $v, v_1, v_2 \in V$ und $\alpha, \alpha_1, \alpha_2 \in K$:

- (v) $(\alpha_1 \cdot (\alpha_2 \cdot v)) = (\alpha_1 \alpha_2) \cdot v$
- (vi) $1 \cdot v = v$
- (vii) $(\alpha_1 + \alpha_2) \cdot v = \alpha_1 \cdot v +_V \alpha_2 \cdot v$
- (viii) $\alpha \cdot (v_1 +_V v_2) = \alpha \cdot v_1 +_V \alpha \cdot v_2$

Hierbei ist $+$ die Addition in K , das Malzeichen ist wie immer unterdrückt. Die Elemente von V heissen *Vektoren*, die Elemente von K *Skalare*. $+_V$ heisst *Vektoraddition* und $\cdot$ *Skalarmultiplikation*. Wegen der Verträglichkeit dieser Operationen schreiben wir auch $+$ statt $+_V$, und lassen $\cdot$ normalerweise weg. 0_V heisst *Nullvektor*.

³⁰Zum Beispiel, ein Körper K oder K -Vektorraum V , siehe Def. 9.1

³¹Der als supp abgekürzte ‘support’ kann natürlich nicht mit dem als $\sup$ symbolisierten ‘Supremum’ aus Def. 6.13 verwechselt werden.

³²“Operation” ist synonym mit einer Verknüpfung der Form $A \times B \rightarrow B$ zwischen im Allgemeinen verschiedenen Mengen, die man sich als “Wirkung von A auf B ” vorstellt.

§ 9. VEKTORRÄUME

Wir vereinbaren vorsorglich auch, dass wir Skalarmultiplikation auch “von rechts” schreiben können, d.h. für $\alpha \in K$, $v \in V$ ist $v \cdot \alpha = v\alpha = \alpha \cdot v = \alpha v \in V$.

Beispiele 9.2. 1. Die triviale abelsche Gruppe $\{0_V\}$ mit der einzig möglichen Skalarmultiplikation $\alpha \cdot 0_V = 0_V$ für alle $\alpha \in K$ ist ein K -Vektorraum.

2. $(K, +)$ mit der Operation $K \times K \ni (\alpha, v) \mapsto \alpha v$ ist ein K -Vektorraum.

3. Für jedes $m \in \mathbb{N}_0$ ist das m -fache kartesische Produkt der auch als *Zeilenvektoren* geschriebenen m -Tupel, $K^m = K \times \dots \times K \ni v = (v_1, \dots, v_m) = (v_1 \ v_2 \ \dots \ v_m)$ ($v_i \in K$ für alle $i = 1, \dots, m$) mit den komponentenweise definierten Operationen³³

$$\begin{aligned} v_1 +_{K^m} v_2 &:= (v_{1,1} + v_{2,1}, \dots, v_{1,m} + v_{2,m}) \\ \alpha \cdot v &:= (\alpha v_1, \dots, \alpha v_m) \end{aligned} \quad (9.1)$$

ein K -Vektorraum. Davon *a priori* nicht zu unterscheiden ist für $n \in \mathbb{N}_0$ die Menge

der *Spaltenvektoren* mit Einträgen in K , $K^n = K \times \dots \times K \ni v = \begin{pmatrix} v^1 \\ \vdots \\ v^n \end{pmatrix}$ mit

$$\begin{aligned} v_1 +_{K^n} v_2 &:= \begin{pmatrix} v_1^1 + v_2^1 \\ \vdots \\ v_1^n + v_2^n \end{pmatrix} & \alpha \cdot v &:= \begin{pmatrix} \alpha v^1 \\ \vdots \\ \alpha v^n \end{pmatrix} \end{aligned} \quad (9.2)$$

Zeilen- und Spaltenvektoren spezialisieren sich für $m = n = 1$ auf 2. und für $m = n = 0$ auf 1.. Der geometrisch-physikalische Unterschied zwischen Zeilen- und Spaltenvektoren ergibt sich (später) aus der *Beziehung* zwischen den beiden Varianten, siehe § 12. Für allgemeines n und $K = \mathbb{R}$ erhält man natürlich Def. 8.1.

4. Für $m, n \in \mathbb{N}_0$ ist der Raum der *Matrizen mit n Zeilen und m Spalten und Einträgen in K*

$$\text{Mat}_{n \times m}(K) = \left\{ A = \begin{pmatrix} a_1^1 & a_2^1 & \dots & a_m^1 \\ a_1^2 & a_2^2 & \dots & a_m^2 \\ \vdots & \vdots & \ddots & \vdots \\ a_1^n & a_2^n & \dots & a_m^n \end{pmatrix} \mid \begin{array}{l} a_j^i \in K \text{ für } i = 1, \dots, n \\ \text{und } j = 1, \dots, m \end{array} \right\} \quad (9.3)$$

und ebenfalls komponentenweise definierter Addition und Skalarmultiplikation ein K -Vektorraum. Mengentheoretisch besteht $\text{Mat}_{n \times m}(K)$ aus durch das kartesische Produkt $\{1, \dots, n\} \times \{1, \dots, m\}$ indizierte Tupel $A = (a_j^i)_{\substack{i=1, \dots, n \\ j=1, \dots, m}}$ von Körperelementen. Die Hoch-/Tiefstellung der Indizes zeigt den Zweck an, für den wir Matrizen verwenden wollen, nämlich zur Darstellung von linearen Abbildungen, siehe § 11.³⁴

5. Der Körper $\mathbb{C}$ der komplexen Zahlen ist ein Vektorraum über $\mathbb{R}$. Hierbei ist die Skalarmultiplikation definiert über die Einbettung (7.4), gefolgt von Multiplikation

³³Führe doch schon besser hier eine gesonderte Notation für Zeilenvektoren ein.

³⁴Es gibt hierfür in der community eine kaum überschaubare Vielfalt von Konventionen. Auch wegen der Begrenztheit des Alphabets ist es allerdings unmöglich, von Anfang bis Ende der Darstellung ein einzelnes davon durchzuhalten.

in $\mathbb{C}$. Die Identifikation von $\mathbb{C}$ mit $\mathbb{R} \times \mathbb{R}$ entlarvt dieses Beispiel dann als den Spezialfall von 3. für $n = 2$ und $K = \mathbb{R}$.

6. Der Körper $\mathbb{R}$ der reellen Zahlen ist ein Vektorraum über $\mathbb{Q}$.

7. Die Menge der Polynome $K[x]$ in einer Variablen (s. Def. 5.8) ist ein K -Vektorraum.

8. Ist X eine beliebige Menge, so ist die Menge der Abbildungen $\mathcal{F}(X, K) = \{f : X \rightarrow K\} = \text{Abb}(X, K)$ mit den *punktweise* definierten Operationen

$$\begin{aligned}(f_1 +_{\mathcal{F}(X, K)} f_2)(x) &:= f_1(x) + f_2(x) \\ (\alpha \cdot f)(x) &:= \alpha f(x)\end{aligned}\tag{9.4}$$

ein K -Vektorraum.

Übungsaufgabe. Man überprüfe die oben angegebenen Vektorraumaxiome in all diesen Beispielen. Beim 8. ist zu beachten, dass im Falle $X = \emptyset$ gemäss der mengentheoretischen Definition 4.1 $\text{Abb}(\emptyset, K) = \{\emptyset\}$ *nicht* leer ist, und ihr einziges Element zum Nullvektor erklärt werden muss (und kann). Es gilt also in diesem Fall $\mathcal{F}(\emptyset, K) = \{0_{\mathcal{F}}\}$.

Weitere aus den Axiomen folgende Rechenregeln:

Lemma 9.3. Für alle $v \in V$, $\alpha \in K$ gilt

- (i) $0 \cdot v = 0_V$;
- (ii) $(-\alpha) \cdot v = -(\alpha \cdot v)$, insbesondere $(-1) \cdot v = -v$;
- (iii) $\alpha \cdot 0_V = 0_V$

Beweis. (i) Unter Verwendung nur der Axiome (i)–(viii) aus Def. 9.1: $0 \cdot v = 0 \cdot v + 0_V = 0 \cdot v + v - v = 0 \cdot v + 1 \cdot v - v = (0 + 1) \cdot v - v = v - v = 0_V$.

(ii) Ebenso: $(-\alpha) \cdot v + \alpha \cdot v = ((-\alpha) + \alpha) \cdot v = 0 \cdot v = 0_V$ wegen (i). Daraus folgt $(-\alpha) \cdot v = -(\alpha \cdot v)$.

(iii) Wegen $0_V + 0_V = 0_V$ ist $-0_V = 0_V$. Dann folgt unter Benutzung von (ii) und (i): $\alpha \cdot 0_V = \alpha \cdot (0_V - 0_V) = \alpha \cdot 0_V + \alpha \cdot (-1) \cdot 0_V = (\alpha - \alpha) \cdot 0_V = 0 \cdot 0_V = 0_V$. $\square$

In der Folge ist es üblich, den Nullvektor 0_V auch als 0 zu schreiben. Dies führt nur selten zu Verwirrung. Beachte allerdings, dass eine “Multiplikation von Vektoren” im allgemeinen Zusammenhang *nicht* erklärt ist (zur Verallgemeinerung des euklidischen inneren Produkts aus § 8 kommen wir erst in HöMa 2). Zur weiteren Untersuchung der Struktur von Vektorräumen halten wir von nun an den Körper K stets fest.

Definition 9.4. Sei V ein K -Vektorraum. Eine nicht-leere Teilmenge $U \subset V$ heisst *Untervektorraum* (häufig auch nur Unterraum, manchmal noch linearer Teilraum, abgekürzt: UVR), falls U unter Vektoraddition und Skalarmultiplikation abgeschlossen ist, d.h. für alle $u, u_1, u_2 \in U$, $\alpha \in K$ gilt $u_1 + u_2 \in U$ und $\alpha \cdot u \in U$.

Lemma 9.5. Jeder Untervektorraum U eines Vektorraums V ist mit den “ererbten” Verknüpfungen (d.h. der Einschränkung von $+_V$ und $\cdot$ auf $U \times U$ bzw. $K \times U$) selbst wieder ein K -Vektorraum.

§9. VEKTORRÄUME

Beweis. Die nicht-triviale einzusehenden Bedingungen sind die Existenz des Nullvektors und des inversen Elements. Letzteres folgt wegen der Abgeschlossenheit unter Skalarmultiplikation aus Lemma 9.3 (ii): Für jedes $u \in U$ ist auch $-u = (-1) \cdot u \in U$. Weil U nicht leer ist, existiert aber auch mindestens ein $u \in U$ und damit folgt wegen der Abgeschlossenheit unter Addition, dass $0_V = u + (-u) \in U$. (Im Falle $U = \{0_V\}$ muss $u = 0_V$ gewählt werden, das ändert aber nichts am Argument.) $\square$

Übungsaufgabe. Eine Teilmenge $U \subset V$ ist genau dann ein linearer Unterraum, wenn (i) $U \neq \emptyset$ und (ii) $\forall \alpha \in K, u_1, u_2 \in U : u_1 + \alpha u_2 \in U$.

Beispiele 9.6. 1. Für jeden Vektorraum V sind $\{0_V\}$ und V “trivialerweise” Unterräume von V .

2. Für jeden Vektor $v \in V$ ist $K \cdot v = \{\alpha \cdot v \mid \alpha \in K\}$ ein Unterraum von V . Für $v \neq 0$ heisst $K \cdot v$ die durch v erzeugte Gerade (durch den Ursprung).

3. Eine Menge der Form $\{v_0 + \alpha v_1 \mid \alpha \in K\}$, die sich geometrisch als Gerade durch v_0 interpretieren lässt, ist (falls $v_0 \notin K \cdot v_1$) *kein* linearer Unterraum von V , sondern lediglich ein *affiner Raum über $K \cdot v_1$* (vgl. Def. 8.2).

4. Für jedes Paar von Vektoren $v_1, v_2 \in V$ ist

$$E := K \cdot v_1 + K \cdot v_2 = \{\alpha_1 v_1 + \alpha_2 v_2 \mid \alpha_i \in K\} \quad (9.5)$$

ein Unterraum von V . Für $v_1 \neq 0_V$ und $v_2 \notin K v_1$ (d.h. v_1 und v_2 sind nicht “kollinear”, dies impliziert insbesondere auch $v_2 \neq 0_V$) heisst E die von v_1 und v_2 erzeugte Ebene (durch den Ursprung).

5. Ist I eine beliebige Indexmenge und $(U_i)_{i \in I}$ eine I -indizierte Familie von Unterräumen eines festen Vektorraums V , so ist auch

$$\bigcap_{i \in I} U_i = \{u \in V \mid u \in U_i \text{ für alle } i \in I\} \quad (9.6)$$

ein Unterraum von V .

6. Ist $(U_1, \dots, U_n)$ eine endliche Familie von Unterräumen von V , so ist auch

$$U_1 + U_2 + \dots + U_n := \{u_1 + u_2 + \dots + u_n \mid u_i \in U_i \text{ für } i = 1, 2, \dots, n\} \quad (9.7)$$

ein Unterraum von V . Beachte, dass die mengentheoretische Vereinigung kein Unterraum sein muss, da sie nicht notwendig unter Addition abgeschlossen ist.

7. Die Menge $K[x]_d$ der Polynome von Grad *kleiner oder gleich* d ist ein Unterraum des Polynomrings $K[x]$, aufgefasst als K -Vektorraum. Die Menge der Polynome mit fixiertem Grad jedoch *nicht*.

8. Ist X eine beliebige Menge, so ist die Menge der *Funktionen mit endlichem Träger* $\mathcal{F}^{\text{fin}}(X, K) = \{f : X \rightarrow K \mid |\text{supp}(f)| < \infty\}$ ein Untervektorraum des Vektorraums $\mathcal{F}(X, K)$ aller Funktionen auf X .

Übungsaufgabe. Man überprüfe die Bedingungen von Def. 9.4 in all diesen Beispielen.

Zu den wesentlichen Aufgaben einer “Strukturtheorie” von Vektorräumen gehört die Klärung der folgenden, teils durch die obigen Beispiele motivierten Fragen: Wann und wie lassen sich gegebene Vektoren zu anderen zusammensetzen, wie beschreiben wir Unterräume allgemein, und inwiefern lässt sich ein gegebener Vektorraum wieder aus seinen Unterräumen zusammensetzen?

Beispiel 9.7. Ist etwa im vierten Beispiel von 9.6 (unter der Voraussetzung, dass $v_1 \neq 0$ und $v_2 \notin K v_1$; dies impliziert insbesondere, dass $v_2 \neq 0$) $u \in E$ ein beliebiger Vektor, so existieren (per Definition) $\alpha_1, \alpha_2 \in K$ so, dass $u = \alpha_1 v_1 + \alpha_2 v_2$. Sind $\beta_1, \beta_2 \in K$ weitere (eventuell andere) Skalare mit $u = \beta_1 v_1 + \beta_2 v_2$, so folgt aus den Rechenregeln zunächst

$$(\alpha_1 - \beta_1)v_1 = -(\alpha_2 - \beta_2)v_2 \quad (9.8)$$

Ist $\alpha_2 - \beta_2 \neq 0$, so folgt

$$v_2 = -\frac{\alpha_1 - \beta_1}{\alpha_2 - \beta_2} v_1 \quad (9.9)$$

d.h. $v_2 \in K \cdot v_1$. Unter der Voraussetzung $v_2 \notin K \cdot v_1$ muss also $\alpha_2 - \beta_2 = 0$ sein. Wäre dann noch $\alpha_1 - \beta_1 \neq 0$, so folgte aus Gl. (9.8) $v_1 = (\alpha_1 - \beta_1)^{-1} 0 = 0$, was wiederum der Voraussetzung $v_1 \neq 0$ widerspräche. Es gilt also $(\beta_1, \beta_2) = (\alpha_1, \alpha_2)$, d.h. die Darstellung von u als Summe von v_1, v_2 ist eindeutig. Wir sagen hierzu (bald): Das Paar (v_1, v_2) ist eine *Basis* von E , siehe Def. 9.15.

Setzen wir andererseits $\tilde{v}_1 = v_1 + v_2$, $\tilde{v}_2 = v_1 - v_2$, so gilt (unter der Voraussetzung, dass $2 \neq 0 \in K$),

$$v_1 = \frac{1}{2}(\tilde{v}_1 + \tilde{v}_2), \quad v_2 = \frac{1}{2}(\tilde{v}_1 - \tilde{v}_2) \quad (9.10)$$

und daher

$$u = \frac{\alpha_1}{2}(\tilde{v}_1 + \tilde{v}_2) + \frac{\alpha_2}{2}(\tilde{v}_1 - \tilde{v}_2) = \tilde{\alpha}_1 \tilde{v}_1 + \tilde{\alpha}_2 \tilde{v}_2 \quad (9.11)$$

wo $\tilde{\alpha}_1 = \frac{1}{2}(\alpha_1 + \alpha_2)$, $\tilde{\alpha}_2 = \frac{1}{2}(\alpha_1 - \alpha_2)$. Die Zerlegung von u bezüglich $\tilde{v}_1, \tilde{v}_2$ "sieht also anders aus". Sie ist aber auch eindeutig, denn es gilt $\tilde{v}_1 \neq 0$ und $\tilde{v}_2 \notin K \cdot \tilde{v}_1$. In Matrix-Schreibweise

$$\begin{aligned} (\tilde{v}_1, \tilde{v}_2) &= (v_1, v_2) \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} & (v_1, v_2) &= (\tilde{v}_1, \tilde{v}_2) \begin{pmatrix} \frac{1}{2} & \frac{1}{2} \\ \frac{1}{2} & -\frac{1}{2} \end{pmatrix} \\ \begin{pmatrix} \alpha_1 \\ \alpha_2 \end{pmatrix} &= \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \begin{pmatrix} \tilde{\alpha}_1 \\ \tilde{\alpha}_2 \end{pmatrix} & \begin{pmatrix} \tilde{\alpha}_1 \\ \tilde{\alpha}_2 \end{pmatrix} &= \begin{pmatrix} \frac{1}{2} & \frac{1}{2} \\ \frac{1}{2} & -\frac{1}{2} \end{pmatrix} \begin{pmatrix} \alpha_1 \\ \alpha_2 \end{pmatrix} \end{aligned} \quad (9.12)$$

Was ist hier los, und was sind die allgemeinen Prinzipien?

Definition 9.8. Sei V ein K -Vektorraum.

(i) Eine *Linearkombination* von endlich vielen (nicht notwendiger Weise verschiedenen!) Vektoren $v_1, v_2, \dots, v_n \in V$ ist ein Vektor der Form

$$\sum_{i=1}^n \alpha_i v_i = \alpha_1 v_1 + \dots + \alpha_n v_n, \quad \text{mit } \alpha_i \in K \text{ für } i = 1, \dots, n \quad (9.13)$$

Hierbei heißen die $\alpha_i \in K$ die *Koeffizienten* der Linearkombination.³⁵ Für gegebene $v_1, v_2, \dots, v_n$ ist die Menge ihrer Linearkombinationen ein Untervektorraum von V . Man nennt diesen Untervektorraum die *lineare Hülle* des Systems oder der Familie (ggfs. auch der Menge) $(v_1, v_2, \dots, v_n)$, und schreibt $\text{span}(v_1, v_2, \dots, v_n)$ (manchmal

³⁵Diese Sprechweise ist (ganz) leicht irreführend, denn die α_i "gehören" der Linearkombination nicht als Element von V , sondern als "Bildungsprozess".

§ 9. VEKTORRÄUME

auch $\text{span}_K(v_1, v_2, \dots, v_n)$, wenn der zugrundeliegende Körper nicht aus dem Kontext völlig klar ist).

(ii) Ist F eine (nicht notwendiger Weise endliche) Menge ($F \subset V$) oder I -indizierte Familie ($F = (v_i)_{i \in I}$) von (nicht notwendiger Weise verschiedenen) Vektoren in V , so heisst $v \in V$ *Linearkombination* von F , falls v Linearkombination einer endlichen Teilmenge oder endlichen Teilfamilie von F ist.³⁶ Die Menge dieser Linearkombinationen ist wieder ein UVR und heisst auch *lineare Hülle* von F , geschrieben $\text{span}(F)$. Andere Sprechweise: $\text{span}(F)$ ist der *von F aufgespannte Unterraum*. Beachte: Im Allgemeinen ist $\text{span}(F)$ verschieden (nämlich eine echte Obermenge) von der linearen Hülle einer jeden festen endlichen Teilmenge/Teilfamilie.

Da wir bei (ii) natürlich auch endliche Mengen oder Indexmengen zulassen, ist (i) einfach nur ein Spezialfall von (ii). Zur Betonung reden wir bei (ii) manchmal auch von *endlichen Linearkombinationen*. Ist F eine Menge, so lassen wir zur Bildung von Linearkombinationen auch Familien von Elementen von F zu. Ist $F = (v_i)_{i \in I}$ eine Familie, so lassen wir auch endliche “Überteilfamilien” zu, also Verknüpfungen $\{1, \dots, n\} \rightarrow I \rightarrow V$, bei denen der erste Pfeil nicht notwendig injektiv ist. Beliebige (endliche) Wiederholung von Vektoren ist also auf jeden Fall erlaubt, auch wenn’s nicht viel bringt.

Proposition 9.9. *Die lineare Hülle einer Familie $(v_i)_{i \in I}$ (oder Menge) von Vektoren $v_i \in V$ ist der “kleinste” lineare Unterraum, der die gesamte Familie/Menge enthält, genauer*

$$\text{span}((v_i)_{i \in I}) = \bigcap_{\substack{U \subset V \text{ UVR} \\ \forall i \in I: v_i \in U}} U \quad (9.14)$$

Beweis. Ist $U \subset V$ ein UVR, der alle v_i enthält, so ist per Def. 9.4 auch jede Linearkombination der v_i in U enthalten. Jede Linearkombination der v_i ist also in jedem UVR auf der rechten Seite enthalten. Es folgt

$$\text{span}((v_i)_{i \in I}) \subset \bigcap_{\substack{U \subset V \text{ UVR} \\ \forall i \in I: v_i \in U}} U \quad (9.15)$$

Andererseits ist die lineare Hülle selbst ein UVR, der alle v_i enthält. Der Schnitt mit den anderen könnte das Ergebnis höchstens verkleinern. Deshalb gilt

$$\bigcap_{\substack{U \subset V \text{ UVR} \\ \forall i \in I: v_i \in U}} U \subset \text{span}((v_i)_{i \in I}) \quad (9.16)$$

Zusammen folgt die Behauptung. □

Wir wollen nun zeigen, dass jeder gegebene Vektorraum durch die Eigenschaften bestimmter Teilmengen komplett festgelegt ist und ökonomisch effizient beschrieben werden kann. Dies geschieht über das Begriffspaar “Erzeugung” (relativ einfach zu verstehen) und “lineare Unabhängigkeit” (etwas subtiler).

³⁶Wir erlauben hier auch die leere Menge, deren (einzige) Linearkombination der Nullvektor ist.

Definition 9.10. Sei V ein K -Vektorraum. Eine Familie $E = (e_i)_{i \in I}$ (oder Menge $E \subset V$) von Vektoren in V heisst *Erzeugendensystem*, falls $\text{span}(E) = V$. V heisst *endlich erzeugt*, falls ein endliches Erzeugendensystem existiert.

Beispiel. Die Familie $\left(\begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 2 \\ 2 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \end{pmatrix}\right)$ ist ein Erzeugendensystem von $\mathbb{R}^2$.

Offensichtlich können wir zu einem gegebenen Erzeugendensystem beliebige Vektoren hinzunehmen (inklusive einzelne Vektoren wiederholen), ohne die Erzeugendeneigenschaft zu verlieren. Endlich erzeugte Vektorräume können auch unendliche Erzeugendensysteme haben. Der gesamte Vektorraum V ist stets Erzeugendensystem seiner selbst.

Definition 9.11. Sei V ein K -Vektorraum. Eine Teilmenge $F \subset V$ heisst *linear unabhängig*, wenn der Nullvektor $0_V \in V$ sich nur als “triviale” Linearkombination von F darstellen lässt. Hiermit ist gemeint: Für jedes endliche System $v_1, \dots, v_n \in F$ und $\alpha_1, \dots, \alpha_n \in K$ gilt

$$\underbrace{(v_i \neq v_j \text{ für alle } i \neq j)}_{\text{Man sagt hierzu "Die } v_i \text{ sind paarweise verschieden".}} \wedge \left(\sum_{i=1}^n \alpha_i v_i = 0 \right) \Rightarrow \alpha_i = 0 \text{ für alle } i = 1, \dots, n \quad (9.17)$$

Beispiel. Die Menge $\left\{v_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}, v_2 = \begin{pmatrix} 1 \\ -1 \end{pmatrix}\right\} \subset \mathbb{R}^2$ ist linear unabhängig: Aus $\alpha_1 v_1 + \alpha_2 v_2 = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$ folgt $\alpha_1 + \alpha_2 = \alpha_1 - \alpha_2 = 0$ und daraus $\alpha_1 = \alpha_2 = 0$.

Offensichtlich ist jede Teilmenge einer linear unabhängigen Menge wieder linear unabhängig. Die leere Menge ist als linear unabhängig einzusehen. Wenn $0_V \in F$, so ist F nicht linear unabhängig.³⁷

Man beachte, dass lineare Unabhängigkeit eine Aussage über die Menge F als Ganzes ist, und nicht über das Verhältnis von Vektoren in der Familie untereinander. Man betrachtet daher den Satz “Die Vektoren von F sind voneinander linear unabhängig” nicht als Aussage im Sinne der Vereinbarung 1.1. Weit verbreitet und zulässig hingegen ist zu sagen “Ein Vektor v ist von einer Menge oder Familie F linear abhängig”, falls “ $v \in \text{span}(F)$ ”.

Lemma 9.12. Sei $F \subset V$ linear unabhängig. Dann ist $v \in V$ genau dann von F linear abhängig, falls entweder $v \in F$ oder $F \cup \{v\}$ nicht linear unabhängig ist.

Beweis. “ $\Rightarrow$ ”: Ist $v \in F$, so ist trivialerweise $v \in \text{span}(F)$. Ist $v \notin F$ und $F \cup \{v\}$ nicht linear unabhängig, so existiert eine nicht-triviale Darstellung des Nullvektors $\sum_{i=1}^n \alpha_i v_i + \alpha v = 0$, mit $v_i \in F$ paarweise verschieden und $\alpha_i, \alpha \in K$ nicht alle 0. Dabei kann α nicht 0 sein, denn sonst wäre mindestens ein $\alpha_i \neq 0$ und bereits

³⁷Definiert man für Familien den Begriff der linearen Unabhängigkeit über (endliche, injektive) Unterfamilien (d.h. $(v_i)_{i \in I}$ ist linear unabhängig, falls für alle endlichen $J \subset I$ aus $\sum_{j \in J} \alpha_j v_j = 0$ folgt, dass $\alpha_j = 0$ für alle $j \in J$), so ist eine linear unabhängige Familie notwendigerweise injektiv und kann (ggfs. bis auf Anordnung) mit ihrem Bild als Teilmenge von V identifiziert werden.

§ 9. VEKTORRÄUME

$\sum_{i=1}^n \alpha_i v_i = 0$ eine nicht-triviale Darstellung des Nullvektors, im Widerspruch zur linearen Unabhängigkeit von F . Dann ist aber

$$v = -\frac{1}{\alpha} \sum_{i=1}^n \alpha_i v_i \quad (9.18)$$

d.h. $v \in \text{span}(F)$ ist von F linear abhängig.

“ $\Leftarrow$ ”: Ist $v \in \text{span}(F)$, aber $v \notin F$, so existiert eine Darstellung $v = \sum_{i=1}^n \alpha_i v_i$ mit $v_i \in F$ paarweise verschieden, und dann ist

$$v - \sum_{i=1}^n \alpha_i v_i = 0 \quad (9.19)$$

eine nicht-triviale Darstellung des Nullvektors, d.h. $F \cup \{v\}$ ist nicht linear unabhängig. $\square$

Lemma 9.13. *$F \subset V$ ist genau dann linear unabhängig, wenn für jeden Vektor $v \in \text{span}(F)$ die Darstellung $v = \sum_{i=1}^n \alpha_i v_i$ als Linearkombination von paarweise verschiedenen $v_1, \dots, v_n \in F$ eindeutig ist bis auf Vertauschen der v_i untereinander und Hinzufügen oder Entfernen von Summanden mit Koeffizient 0.*

Beweisskizze. Ist die Darstellung für jeden Vektor eindeutig, so insbesondere auch für $v = 0$, d.h. F ist linear unabhängig. Ist umgekehrt F linear unabhängig und sind $v = \sum_{i=1}^n \alpha_i v_i = \sum_{j=1}^m \beta_j w_j$ zwei (a priori verschiedene) Darstellungen von $v \in \text{span}(F)$ (mit $v_i, w_j \in F$ jeweils paarweise verschieden und $\alpha_i, \beta_j \in K$), so können wir zunächst durch Auffüllen mit 0 und Anpassen der Reihenfolge erreichen, dass $m = n$ und $w_i = v_i$ für alle i (aber immer noch paarweise verschieden). Dann ist aber

$$\sum_{i=1}^n (\alpha_i - \beta_i) v_i = v - v = 0 \quad (9.20)$$

eine Darstellung des Nullvektors, und wegen der linearen Unabhängigkeit von F muss $\alpha_i - \beta_i = 0$ sein für alle $i = 1, \dots, n$. $\square$

Die Vergrößerung einer linear unabhängigen Menge ist nicht unbedingt mehr linear unabhängig, während die Verkleinerung eines Erzeugendensystems nicht unbedingt ein Erzeugendensystem mehr ist. Auf der Suche nach dem Mittelweg ist die zentrale Aussage der sog. “Steinitzsche Austauschsatz”, den wir allerdings nur für endlich erzeugte Vektorräume formulieren und beweisen.

Proposition 9.14. *Sei V ein endlich erzeugter K -Vektorraum, und $E = (e_1, \dots, e_k)$ ein endliches Erzeugendensystem von V . Ist $B \subset V$ linear unabhängig, so ist B endlich und $\#B \leq k$.*

Beweis. Für $B = \emptyset$ ist nichts zu zeigen. Andernfalls sei $b_1 \in B$, notwendigerweise $b_1 \neq 0$. Da E V erzeugt, existieren $\alpha_1^j \in K$ für $j = 1, \dots, k$ so, dass

$$b_1 = \sum_{j=1}^k \alpha_1^j e_j \quad (9.21)$$

Mindestens eines der α_1^j ist nicht 0 (sonst wäre $b_1 = 0$). Ohne Einschränkung (bzw. nach Umnummerieren der e_j) nehmen wir an, dass $\alpha_1^1 \neq 0$. Dann ist

$$e_1 = \frac{1}{\alpha_1^1} \left(b_1 - \sum_{j=2}^k \alpha_1^j e_j \right) \quad (9.22)$$

und jede Linearkombination von $(e_1, e_2, \dots, e_k)$ lässt sich auch als Linearkombination von $(b_1, e_2, \dots, e_k)$ schreiben:

$$\sum_{j=1}^k \beta^j e_j = \beta^1 \left(\frac{1}{\alpha_1^1} \left(b_1 - \sum_{j=2}^k \alpha_1^j e_j \right) \right) + \sum_{j=2}^k \beta^j e_j = \frac{\beta^1}{\alpha_1^1} b_1 + \sum_{j=2}^k \left(\beta^j - \frac{\beta^1 \alpha_1^j}{\alpha_1^1} \right) e_j \quad (9.23)$$

$E_1 := (b_1, e_2, \dots, e_k)$ ist also ebenfalls ein Erzeugendensystem von V .

Existiert ein weiteres $b_2 \in B \setminus \{b_1\}$, notwendigerweise $b_2 \neq 0$, so lässt sich dieses als Linearkombination von E_1 schreiben,

$$b_2 = \alpha_2^1 b_1 + \sum_{j=2}^k \alpha_2^j e_j \quad (9.24)$$

Wäre $\alpha_2^j = 0$ für alle $j = 2, \dots, k$, so wäre $\{b_1, b_2\} \subset B$ nicht linear unabhängig, was gemäss Bemerkung unter Def. 9.11 nicht sein kann. Ist ohne Einschränkung $\alpha_2^2 \neq 0$, so können wir analog zu oben e_2 durch b_2 ersetzen und erhalten das Erzeugendensystem $E_2 := (b_1, b_2, e_3, \dots, e_k)$. Usw.

Bei der Fortsetzung dieses Verfahrens muss B spätestens nach k Schritten ausgeschöpft sein: Existierte (falls das Verfahren nicht schon früher abbricht) nach der Konstruktion des Erzeugendensystems $E_k = (b_1, \dots, b_k)$ noch ein $b_{k+1} \in B \setminus E_k$, so wäre b_{k+1} Linearkombination von E_k und damit $E_k \cup \{b_{k+1}\}$ nicht linear unabhängig (s. Lemma 9.12). Die Mächtigkeit von B ist also nicht grösser als k . $\square$

Wir sind nun bereit für die zentrale Definition.

Definition 9.15. Sei V ein K -Vektorraum. Eine Teilmenge $B \subset V$ heisst *Basis von V* , wenn gilt:

- (i) B ist Erzeugendensystem von V , d.h. $V = \text{span}(B)$
- (ii) B ist linear unabhängig.

Eine Familie $B = (b_i)_{i \in I}$ heisst Basis, wenn sie injektiv ist und ihr Bild als Teilmenge von V eine Basis im obigen Sinne. Ist $I = \mathbb{N}$ oder von der Form $\{1, 2, \dots, n\}$, so sagen wir auch *geordnete Basis*.

Beispiele 9.16. 1. Die Menge der Vektoren $\{e_1, \dots, e_n\}$, wo

$$e_i = \begin{pmatrix} 0 \\ \vdots \\ 1 \\ \vdots \\ 0 \end{pmatrix} \quad i\text{-te Zeile} \quad (9.25)$$

§ 9. VEKTORRÄUME

ist eine Basis des K^n , genannt die *Standardbasis*: Jedes $v \in K^n$ lässt sich offensichtlich als Linearkombination der e_i schreiben und aus $\sum_{i=1}^n \alpha^i e_i = 0$ folgt sofort $\alpha^i = 0$ für alle i .

2. Die Vektoren $\tilde{e}_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ und $\tilde{e}_2 = \begin{pmatrix} 1 \\ -1 \end{pmatrix}$ bilden eine Basis des K^2 (falls $2 \neq 0 \in K$)

(Übungsaufgabe, vgl. auch Beispiel 9.7).

3. Die Familie von Matrizen (E_i^j) mit $i = 1, \dots, n$, $j = 1, \dots, m$ und

$$E_i^j := \begin{pmatrix} 0 & \cdots & 0 & \cdots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \cdots & 1 & \cdots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \cdots & 0 & \cdots & 0 \end{pmatrix} \begin{matrix} j\text{-te Spalte} \\ \\ i\text{-te Zeile} \end{matrix} \quad (9.26)$$

ist eine Basis von $\text{Mat}_{n \times m}(K)$.

4. Ist X eine beliebige Menge, so ist die Familie von Funktionen $(\delta_x)_{x \in X}$

$$\delta_x(y) = \begin{cases} 1 & \text{falls } y = x \\ 0 & \text{sonst} \end{cases} \quad (9.27)$$

eine Basis des Vektorraums $\mathcal{F}^{\text{fin}}(X, K)$ der Funktionen mit endlichem Träger. (vgl. Beispiel 8 in 9.6). (Sie ist allerdings *keine Basis* des Raums $\mathcal{F}(X, K)$ aller Funktionen: Eine Funktion mit unendlichem Träger lässt sich nicht als endliche Linearkombination der δ_x darstellen.)

Dass Basen (auch ohne weitere Endlichkeitsvoraussetzungen) das “Goldene Mittel” zwischen Erzeugung und Erzeugendensystemen bilden, formulieren wir wie folgt.

Lemma 9.17. *Die folgenden Aussagen sind äquivalent.*

(i) B ist eine Basis von V .

(ii) B ist minimales Erzeugendensystem. D.h. B ist Erzeugendensystem, aber jede echte Teilmenge ist kein Erzeugendensystem.

(iii) B ist eine maximal linear unabhängige Teilmenge. D.h. B ist linear unabhängig, aber jede echte Obermenge ist nicht linear unabhängig.

Beweis. (i) $\Rightarrow$ (ii),(iii): Ist B eine Basis, und $F \subsetneq B$, so existiert ein $v \in B \setminus F$. v ist nicht Linearkombination von F , sonst wäre B nicht linear unabhängig. Das heisst F ist kein Erzeugendensystem. Ist $B \subsetneq E$, dann existiert ein $v \in E \setminus B$. Da B Erzeugendensystem ist, ist $v \in \text{span}(B)$, d.h. E ist nicht linear unabhängig.

(ii) $\Rightarrow$ (i): Ist B minimales Erzeugendensystem, aber nicht linear unabhängig, so existiert eine nicht-triviale Darstellung $\sum \alpha_i v_i = 0$ des Nullvektors mit $v_i \in B$ mit (say) $\alpha_1 \neq 0$. Dann aber ist $B \setminus \{v_1\} \subsetneq B$ ebenfalls Erzeugendensystem. $\nexists$

(iii) $\Rightarrow$ (i): Ist B maximal linear unabhängige Teilmenge, aber kein Erzeugendensystem, so existiert ein $v \in V \setminus \text{span}(B)$. Dann aber ist $B \cup \{v\} \supsetneq B$ linear unabhängig. $\nexists$ □

Andere Variante:

Lemma 9.18. *$B \subset V$ ist genau dann eine Basis, wenn sich jeder Vektor $v \in V$ auf genau eine Weise als (endliche) Linearkombination von (paarweise verschiedenen) Elementen von B darstellen lässt.*

Beweis. Folgt unmittelbar aus Lemma 9.17 und Lemma 9.13. $\square$

Ausserdem gilt:

Proposition 9.19 (Basisergänzungssatz). *Sei $F \subset V$ linear unabhängig, $E \subset V$ ein Erzeugendensystem von V , und $F \subset E$. Dann existiert eine Basis B von V so, dass $F \subset B \subset E$.*

Beweis. (für endlich erzeugte Vektorräume). Ist F nicht schon selbst Erzeugendensystem von V , dann gilt sogar schon $E \not\subset \text{span}(F)$ (sonst wäre jeder Vektor in der linearen Hülle von E , d.h. also jeder Vektor in v , auch in der linearen Hülle von F). Also existiert ein $e_1 \in E \setminus \text{span}(F)$, so dass $F_1 = F \cup \{e_1\}$ linear unabhängig ist (siehe Lemma 9.12). Ist dies immer noch kein Erzeugendensystem von V , so wiederholen wir die Prozedur. Sie muss nach endlich vielen Schritten abbrechen, denn wegen Prop. 9.14 ist die Anzahl Elemente von linear unabhängigen Mengen beschränkt. (Wobei nicht vorausgesetzt wird, dass E selbst eines der endlichen Erzeugendensysteme von V ist.)

Der allgemeine Nachweis (für nicht endlich-erzeugte Vektorräume) hängt via dem “Zornschen Lemma” ab von dem Auswahlaxiom und ist demnach “nicht konstruktiv”. Bekannte Beispiele, für die die explizite Angabe einer Basis nicht gelingt, sind $\mathbb{R}$ als $\mathbb{Q}$ -Vektorraum oder der Vektorraum $\mathcal{F}(\mathbb{N}, \mathbb{R})$ aller reellen Folgen. $\square$

Für endlich erzeugte Vektorräume halten wir fest:

Theorem 9.20 (Hauptsatz für endlich erzeugte Vektorräume). *Sei V ein endlich erzeugter K -Vektorraum. Dann gilt*

- (i) *V besitzt eine endliche Basis $B = \{b_1, \dots, b_n\}$.*
- (ii) *Sind B und $\tilde{B}$ zwei Basen von V , so gilt $\#B = \#\tilde{B}$.*

Beweis. (i) folgt, ausgehend von $F = \emptyset$, direkt aus Prop. 9.19. Für (ii) benutzen wir Prop. 9.14: Aus der linearen Unabhängigkeit von B und der Erzeugendeneigenschaft von $\tilde{B}$ folgt $\#B \leq \#\tilde{B}$, und umgekehrt $\#\tilde{B} \leq \#B$, zusammen die Behauptung. $\square$

Definition 9.21. Ist V endlich erzeugt, so heisst die gemäss Theorem 9.20 von der Wahl einer Basis B unabhängige natürliche Zahl $\dim V = \#B \in \mathbb{N}$ die *Dimension* von V , und V auch *endlich-dimensional*. Ist der zugrunde liegende Vektorraum nicht aus dem Kontext klar, so schreiben wir auch $\dim_K V$. Ist V nicht endlich erzeugt, so sagt man auch V ist *unendlich-dimensional* und schreibt $\dim V = \infty$.

Beispiele 9.22. 1. $\dim K^n = n$ (siehe 1. Beispiel 9.16).

2. Ist X eine endliche Menge, so ist $\mathcal{F}(X, K) = \mathcal{F}^{\text{fin}}(X, K)$ und endlich-dimensional, mit $\dim \mathcal{F}(X, K) = \#X$ (siehe 4. Beispiel 9.16)

Definition 9.23. Sei V ein endlich erzeugter K -Vektorraum, und $B = (b_1, \dots, b_n)$ eine Basis von V . Dann heissen die Koeffizienten v^i in der Darstellung von $v \in V$ als Linearkombination $v = \sum_{i=1}^n v^i b_i$ von B (welche gemäss Lemma 9.18 eindeutig ist) die (linearen) *Koordinaten* von v bezüglich B .

§ 9. VEKTORRÄUME

Definition 9.24. Ist X ein affiner Raum über V (s. Def. 8.2 und Lemma 8.3), $x_0 \in X$, und B eine Basis von V , so heissen für $x \in X$ die Koordinaten von $\delta(x_0, x) \in V$ bezüglich B *lineare Koordinaten* von x bezüglich B und x_0 .

Bemerkung. Wir beschränken uns im weiteren Verlauf der HöMa 1 auf endlich-dimensionale (bzw. synonym endlich erzeugte) Vektorräume.³⁸ Im Kontext unendlich-dimensionaler (aber normierter, s. § 15) Vektorräume, wie sie in der HöMa 2 und 3 behandelt werden, erlaubt man auch bei Erzeugung und linearer Unabhängigkeit unendliche Linearkombinationen. Die resultierenden Basen heissen dann manchmal “Schauder”, die gewöhnlichen algebraischen auch “Hamel”.

Wir halten noch einige einfache aber sehr nützliche Folgerungen aus den obigen Erkenntnissen fest.

Korollar 9.25. Sei V ein endlich erzeugter K -Vektorraum mit $\dim_K(V) = n$.

(i) Ist $F \subset V$ linear unabhängig, so gilt $\#F \leq n$ und $\#F = n$ genau dann, wenn F eine Basis ist.

(ii) Ist $E = (e_i)_{i \in I}$ ein Erzeugendensystem, so gilt $\#I \geq n$ und $\#I = n$ genau dann, wenn E eine Basis ist.

Beweis. Die Aussagen folgen aus den Charakterisierungen von Basen in Lemma 9.17 und der Invarianz der Basislänge, Theorem 9.20 (ii). $\square$

Korollar 9.26. Ist V ein endlich erzeugter Vektorraum, und $U \subset V$ ein linearer Unterraum, so ist U ebenfalls endlich erzeugt und es gilt $\dim U \leq \dim V$. Gleichheit ist äquivalent zu $U = V$. Jede Basis von U lässt sich zu einer Basis von V ergänzen.

Beweis. Wegen Prop. 9.14 ist die Mächtigkeit linear unabhängiger Teilmengen von U beschränkt (nämlich durch $\dim V$), und wie im Beweis von Prop. 9.19 können wir eine Basis A von U als maximal linear unabhängige Teilmenge explizit konstruieren. $\dim U \leq \dim V$ folgt dann aus Prop. 9.14, die Ergänzzbarkeit von A zu einer Basis von V aus Prop. 9.19. Ist $\dim U = \dim V$, so muss die Ergänzung trivial sein, d.h. A ist bereits eine Basis von V . $\square$

Proposition 9.27. Es seien U und W Unterräume eines endlich-dimensionalen Vektorraums. Dann gilt

$$\dim U + \dim W = \dim(U + W) + \dim(U \cap W) \quad (9.28)$$

(Vgl. Beispiele 5 und 6 in 9.6 für die Definition von $U \cap W$ und $U + W$.)

Beweis. Sei $n := \dim(U \cap W)$ und $\{v_1, \dots, v_n\}$ eine Basis von $U \cap W$. Ist $\dim U = k$ und $\dim W = m$, so existieren gemäss Kor. 9.26 Vektoren $u_{n+1}, \dots, u_k \in U$ und Vektoren $w_{n+1}, \dots, w_m \in W$ so, dass $\{v_1, \dots, v_n, u_{n+1}, \dots, u_k\}$ Basis von U ist und $\{v_1, \dots, v_n, w_{n+1}, \dots, w_m\}$ Basis von W . Die Behauptung, $\dim(U + W) = k + m - n$, folgt dann aus der Wahrheit der Aussage, dass

$$B = \{v_1, \dots, v_n, u_{n+1}, \dots, u_k, w_{n+1}, \dots, w_m\}$$

³⁸In diesem Kapitel betrachten wir nur endliche, ab § 16 auch unendliche Linearkombinationen, aber immer innerhalb eines festen endlich-dimensionalen Vektorraums.

eine Basis von $U + W$ ist. B ist aber offensichtlich Erzeugendensystem, zu zeigen bleibt also lineare Unabhängigkeit: Ist³⁹

$$\sum_{i=1}^n \alpha_i v_i + \sum_{i=n+1}^k \beta_i u_i + \sum_{i=n+1}^m \gamma_i w_i = 0 \quad (9.29)$$

so ist der Vektor

$$X = \sum_{i=1}^n \alpha_i v_i + \sum_{i=n+1}^k \beta_i u_i = - \sum_{i=n+1}^m \gamma_i w_i \quad (9.30)$$

sowohl in U als auch in W enthalten. Er liegt innerhalb von W also im Unterraum $U \cap W \subset W$ und muss sich demnach als Linearkombination nur der v_i darstellen lassen. Daraus folgt $\gamma_i = 0$ für $i = n + 1, \dots, m$. Das heisst aber $X = 0$ auch als Element von U betrachtet und daraus folgt $\beta_i = 0$ für $i = n + 1, \dots, k$ und $\alpha_i = 0$ für $i = 1, \dots, n$. $\square$

§ 10 Basen, Matrizen, Determinanten

Die abstrakten Betrachtungen des letzten § scheinen zunächst wenig darüber auszusagen, welche endlich-dimensionalen Vektorräume es ausser dem K^n noch gibt,⁴⁰ wie sie zu konstruieren wären, bzw. wie man überhaupt, bevor man nach einer Basis sucht, entscheiden soll, ob ein gegebener Vektorraum endlich erzeugt ist oder nicht—Es sei denn (siehe Kor. 9.26) es handelt sich um einen (dann beliebigen) Unterraum eines endlich erzeugten Vektorraums! Wir entdecken hier die entsprechenden Verfeinerungen des Steinitzschen Austauschsatzes (Proposition 9.14), und ihre praktische Anwendung zur Lösung mehrerer grundlegender Rechenaufgaben der linearen Algebra. Für weitere Konstruktionen von und mit endlich-dimensionalen Vektorräumen siehe § 12.

Hierfür benutzen wir Matrizen,⁴¹ die wir glücklicherweise noch gerade auf S. 45 untergebracht hatten, und deren Multiplikation, die wir sträflicherweise nicht mehr am Ende von § 5 erklärt hatten (s. allerdings (9.12)):

Definition 10.1. Sind $l, n, m \in \mathbb{N}_0$, $A = (a_i^h) \in \text{Mat}_{l \times n}(K)$, $B = (b_j^i) \in \text{Mat}_{n \times m}(K)$ so ist das *Produkt* von A mit B die Matrix $C = A \cdot B = AB = (c_j^h) \in \text{Mat}_{l \times m}(K)$ mit Einträgen

$$c_j^h = \sum_{i=1}^n a_i^h b_j^i \quad (10.1)$$

³⁹Wir könnten hier schon feststellen, dass die Elemente von B paarweise verschieden sind. Ansonsten zeigten wir einfaches etwas stärkeres.

⁴⁰Auch der anfangs etwas mysteriöse Raum der Funktionen auf einer endlichen Menge ist für eine gegebene Abzählung $X \simeq \{1, \dots, n\}$ nichts anderes als $\mathcal{F}(X, K) = \{(v_i)_{i=1, \dots, n}\} = K^n$.

⁴¹In diesem Einschub bezeichnen wir Matrizen mit Grossbuchstaben vom Anfang des Alphabets. Im weiteren Verlauf des § werden sie kalligraphisch notiert, um sie von Mengen/Tupel von Vektoren eines abstrakten Vektorraums zu unterscheiden. Auch für sonstige Vorschläge zur Optimierung der Notation bin ich (unter Vorbehalt) prinzipiell empfänglich.

§ 10. BASEN, MATRIZEN, DETERMINANTEN

In Worten: Der (h, j) -te Eintrag von C entsteht aus der h -ten Zeile von A und der j -ten Spalte von B durch Multiplikation ihrer korrespondierenden Einträge und Summation der Ergebnisse.

- Beachte, dass die Multiplikation nur für “passende Matrizen” definiert ist, wenn also die Anzahl Spalten von A gleich der Anzahl Zeilen von B ist.
- Andere Formulierungen: Die h -te Zeile von C (der Länge m) ist eine Linearkombination der Zeilen von B mit den Einträgen der h -ten Zeile von A als Koeffizienten (“Zeile mal Spalte”). Die j -te Spalte von C (der Höhe l) ist eine Linearkombination der Spalten von A mit den Einträgen der j -ten Spalte von B als Koeffizienten (“Spalte mal Zeile”).
- Spezialfall $m = 1$: Multiplikation eines Spaltenvektors mit einer Matrix “von links”.
- Spezialfall $l = 1$: Multiplikation eines Zeilenvektors mit einer Matrix “von rechts”.
- Trivialfälle: Eine von l, m, n ist gleich 0.
- Vorteil der Indexplatzierung: Wir können unter der *Einsteinschen Summenkonvention* das Summenzeichen weglassen und schreiben ab ziemlich bald:

$$c_j^h = a_i^h b_j^i \quad (10.2)$$

Proposition 10.2. (i) Multiplikation von Matrizen (falls definiert) ist assoziativ, d.h. für alle $A \in \text{Mat}_{l \times n}(K)$, $B \in \text{Mat}_{n \times m}$, $C \in \text{Mat}_{m \times o}(K)$ gilt

$$(A \cdot B) \cdot C = A \cdot (B \cdot C) \quad (10.3)$$

(ii) Multiplikation von Matrizen ist distributiv über der Addition und verträglich mit der Skalarmultiplikation, d.h. für alle $A, A_1, A_2 \in \text{Mat}_{l \times n}(K)$, $B, B_1, B_2 \in \text{Mat}_{n \times m}(K)$, $\alpha, \beta \in K$ gilt:

$$\begin{aligned} (\alpha A_1 + A_2) \cdot B &= \alpha A_1 \cdot B + A_2 \cdot B \\ A \cdot (\beta B_1 + B_2) &= \beta A \cdot B_1 + A \cdot B_2 \end{aligned} \quad (10.4)$$

Beweis. durch Rückführung auf Rechenregeln in K . □

Definition 10.3. Wir schreiben für den unitären Ring der quadratischen $n \times n$ Matrizen auch $\text{Mat}_n(K)$. Sein neutrales Element

$$\mathbb{1}_n = I_n = 1_{\text{Mat}_n(K)} = \begin{pmatrix} 1 & 0 & \cdots & 0 \\ 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & 1 \end{pmatrix} \quad (10.5)$$

heißt *Einheitsmatrix*. Die Einheitengruppe (s. Def. 5.6)

$$\text{Mat}_n(K)^\times = \{A \in \text{Mat}_n(K) \mid \exists A^{-1} : A \cdot A^{-1} = A^{-1} \cdot A = I_n\} \quad (10.6)$$

heißt *allgemein lineare Gruppe* vom Grad n über K , geschrieben $GL(n, K)$.

Bemerkungen. · Ausser für $n = 0$ ($\text{Mat}_0(K) = GL(0, K) = \{e\}$) und $n = 1$ ($\text{Mat}_1(K) = K$, $GL(1, K) = K^\times = K \setminus \{0\}$) ist $\text{Mat}_n(K)$ nicht kommutativ.

· Im Unterschied zur Situation bei allgemeinen Ringen (s. Def. 5.6), impliziert für Matrizen die Existenz einer Linksinversen bereits die Existenz der Rechtsinversen (und umgekehrt). Siehe Lemma 10.7.

· Für $n > 1$ sind nicht alle von der Null verschiedenen Matrizen invertierbar, dies im Unterschied zur Körperstruktur von $\mathbb{C}$ mit der expliziten Formel (7.1) für die Invertierung von “Paaren reeller Zahlen”.

Beispiel. $\begin{pmatrix} 1 & 1 \\ -1 & 1 \end{pmatrix} \cdot \begin{pmatrix} \frac{1}{2} & -\frac{1}{2} \\ \frac{1}{2} & \frac{1}{2} \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$. Die Matrix $\begin{pmatrix} 1 & 2 \\ 2 & 4 \end{pmatrix}$ ist nicht invertierbar, siehe Begründung am Ende des §.

Basisfinden

Problemstellung: Gegeben (Spaltenvektoren) $v_1, \dots, v_m \in K^n$ als Erzeugendensystem $E = (v_i)$ eines endlich-dimensionalen Vektorraums $U = \text{span}(E) \subset K^n$, finde eine Basis von U und bestimme dessen Dimension.

Idee: Wir modifizieren E , zum Unterschied von Prop. 9.14 diesmal ohne U zu verändern, solange bis wir sie auf eine offensichtlich maximal linear unabhängige Teilmenge reduzieren können.

Beobachtungen: • Der von E erzeugte Unterraum ändert sich nicht, wenn wir

I. Jede Nennung des Nullvektors aus der Liste streichen: Er trägt zu Linearkombinationen gar nicht bei;

II. Zwei Vektoren vertauschen: Das ist offensichtlich;

III. Einen der Vektoren durch ein nicht-triviales skalare Vielfache ersetzen: Ist $\lambda \in K^\times = K \setminus \{0\}$, und $1 \leq i \leq m$, so gilt

$$\sum_{j=1}^m \alpha_j v_j = \frac{\alpha_i}{\lambda} (\lambda v_i) + \sum_{j \neq i} \alpha_j v_j \quad (10.7)$$

d.h. $(v_1, \dots, \lambda v_i, \dots, v_m)$ ist genauso ein Erzeugendensystem wie E .

IV. Ein Vielfaches eines Vektors zu einem anderen dazuzugaddieren: Ist $1 \leq i_1 < i_2 \leq m$ und $\lambda \in K$ so gilt

$$\sum_{j=1}^m \alpha_j v_j = (\alpha_{i_1} - \lambda \alpha_{i_2}) v_{i_1} + \alpha_{i_2} (v_{i_2} + \lambda v_{i_1}) + \sum_{j \neq i_1, i_2} \alpha_j v_j \quad (10.8)$$

d.h. $(v_1, \dots, v_{i_2} + \lambda v_{i_1}, \dots, v_m)$ ist genauso ein Erzeugendensystem wie E .

• Andererseits ist die lineare Unabhängigkeit unmittelbar einfach einzusehen, wenn die Matrix $\mathcal{E} = (v_j^i)_{\substack{i=1, \dots, n \\ j=1, \dots, m}} \in \text{Mat}_{n \times m}(K)$, die aus der Darstellung $v_j = v_j^i e_i$ der v_j als Linearkombination der Standardbasis des “umgebenden” K^n (m.a.W.: den v_j als Spaltenvektoren) gebildet wird, in sog. Stufenform ist, d.h. formal: Es existieren Zahlen $1 \leq k \leq m$ und $1 \leq i_1 < i_2 < \dots < i_k \leq n$ so dass

(i) $v_j^i = 0$ falls $j > k$ oder falls $j \leq k$ und $i < i_j$

(ii) $v_j^{i_j} \neq 0$

Anschaulich:

$$\mathcal{E} = \left(\begin{array}{cccc|cccc} 0 & 0 & \cdots & 0 & \cdots & 0 & & \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots & & \\ \hline v_1^{i_1} & 0 & \cdots & 0 & \cdots & 0 & & \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots & & \\ \hline v_1^{i_2} & v_2^{i_2} & \cdots & 0 & \cdots & 0 & & \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots & & \\ \hline v_1^{i_k} & v_2^{i_k} & \cdots & v_k^{i_k} & \cdots & 0 & & \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots & & \\ \hline v_1^n & v_2^n & \cdots & v_k^n & \cdots & 0 & & \end{array} \right) \quad (10.9)$$

Begründung: Aus $0 = \sum_{j=1}^m \alpha^j v_j = \sum_{j=1}^k \alpha^j v_j$ folgt durch Betrachten des i_1 -ten Eintrags wegen $v_1^{i_1} \neq 0$, dass $\alpha^1 = 0$, sodann durch Betrachten des i_2 -ten Eintrags $\alpha^2 = 0$, etc. bis $\alpha^k = 0$. Die ersten k Spalten von $\mathcal{E}$ sind also linear unabhängig.

· Die Zahlen $\alpha_j^{i_j} \neq 0$ für $j = 1, \dots, k$ heissen *Pivot-Elemente* der Stufenform.

Lösung: (Gauss-Algorithmus zur Herstellung der Stufenform) Wir nutzen die Transformation II–IV von S. 58, die als Operationen auf Matrizen *elementare Spaltenumformungen* genannt werden.⁴²

1. Schritt: Setze $i_1 = \text{“kleinstes } i, \text{ so dass } \exists j : v_j^{i_1} \neq 0\text{”}$. Tausche entsprechende Spalte mit der ersten.

2. Schritt: Addiere für alle $j > 1$ das $-v_j^{i_1}/v_1^{i_1}$ -Vielfache der ersten Spalte zur j -ten dazu. Sodann hat $\mathcal{E}$ bis einschliesslich zur i_1 -ten Zeile die gewünschte Form.

3. Schritt: Stelle die erste Spalte zur Seite und wiederhole das Verfahren (sinngemäss) für die verbliebene Matrix.

Ergebnis: Die ersten k Spalten der umgeformten Matrix $\mathcal{E}$ sind eine Basis des von den ursprünglichen m Vektoren aufgespannten Unterraums des K^n .

Beispiel:

$$\begin{aligned} & \begin{pmatrix} 0 & 1 & 1 & -1 \\ 0 & 0 & 0 & 0 \\ -2 & -1 & 0 & 1 \\ 2 & 2 & 1 & 0 \\ -2 & 2 & 3 & -1 \end{pmatrix} \xrightarrow{v_1 \leftrightarrow v_3} \begin{pmatrix} 1 & 1 & 0 & -1 \\ 0 & 0 & 0 & 0 \\ 0 & -1 & -2 & 1 \\ 1 & 2 & 2 & 0 \\ 3 & 2 & -2 & -1 \end{pmatrix} \xrightarrow{v_2 \leftarrow v_2 - v_1} \begin{pmatrix} 1 & 0 & 0 & -1 \\ 0 & 0 & 0 & 0 \\ 0 & -1 & -2 & 1 \\ 1 & 1 & 2 & 0 \\ 3 & -1 & -2 & -1 \end{pmatrix} \\ & \xrightarrow{v_4 \leftarrow v_4 + v_1} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & -1 & -2 & 1 \\ 1 & 1 & 2 & 1 \\ 3 & -1 & -2 & 2 \end{pmatrix} \xrightarrow{v_3 \leftarrow v_3 - 2v_2} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & -1 & 0 & 1 \\ 1 & 1 & 0 & 1 \\ 3 & -1 & 0 & 2 \end{pmatrix} \\ & \xrightarrow{v_4 \leftarrow v_4 + v_2} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 1 & 1 & 0 & 2 \\ 3 & -1 & 0 & 1 \end{pmatrix} \xrightarrow{v_3 \leftrightarrow v_4} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 1 & 1 & 2 & 0 \\ 3 & -1 & 1 & 0 \end{pmatrix} \end{aligned} \quad (10.10)$$

⁴²Tatsächlich benötigen wir hier zunächst nur II und IV.

Der von den Spalten der Ausgangsmatrix aufgespannte Unterraum des K^5 hat also Dimension 3 (und zwar auch wenn $2 = 0 \in K$, haha).

Ambiguitäten: Der Gaußsche Algorithmus hängt (in dem als 1. angegebenen Schritt) bereits von Wahlen ab (nämlich welche der Spalten mit $v_j^{i_1} \neq 0$ wir mit der ersten tauschen). Ausserdem können wir aus dem Ergebnis mit weiteren Operationen der Art II, III und IV andere Basen von U konstruieren. Keine davon ist a priori ausgezeichnet.

Nächste Frage: Wie stehen verschiedene Basen allgemein zueinander in Beziehung, und wie sieht die “Menge aller Basen” eines festen Vektorraums genau aus?

Basiswechsel

Definition 10.4. Sei V ein endlich-dimensionaler Vektorraum mit $\dim V = n$, und $B = (b_1, \dots, b_n)$ eine Basis⁴³ von V . Ist $C = (c_1, \dots, c_n)$ eine weitere Basis von V , so heisst die aus den Koeffizienten der Darstellung⁴⁴ der c_j als Linearkombination von B ,

$$c_j = \sum_{i=1}^n \sigma_j^i b_i = b_i \sigma_j^i \quad (10.11)$$

gebildete Matrix,

$$\mathcal{S}_{BC} = (\sigma_j^i)_{i,j=1,\dots,n} \quad (10.12)$$

die *Basiswechselmatrix* von B zu C .⁴⁵

Lemma 10.5. (i) Ist $v \in V$, $v = v^i b_i$ seine Darstellung als Linearkombination von B , und $v = w^j c_j$ seine Darstellung als Linearkombination von C , so gilt

$$v^i = \sigma_j^i w^j \quad (10.13)$$

(ii) Die Basiswechselmatrix von B nach C ist invertierbar, und es gilt

$$(\mathcal{S}_{BC})^{-1} = \mathcal{S}_{CB} \quad (10.14)$$

(iii) Ist $D = (d_1, \dots, d_n)$ eine dritte Basis von V , so gilt

$$\mathcal{S}_{BD} = \mathcal{S}_{BC} \mathcal{S}_{CD} \quad (10.15)$$

In Worten: Die Basiswechselmatrix enthält einerseits die Koordinaten der Elemente der “neuen” Basis bezüglich der “alten” (Gl. (10.11)) und drückt andererseits für beliebige Vektoren deren “alte” Koordinaten durch die “neuen” aus (Gl. (10.13)).⁴⁶ Beachte auch noch einmal, dass die Matrix $\mathcal{S}_{BC}$ zwischen beliebigen Basen eines “abstrakten” Vektorraums vermittelt. Zur Identifikation mit dem K^n kommen wir erst in § 11 wieder zurück.

⁴³Die runden Klammern zeigen an, dass wir uns von nun an für *geordnete Basen* (endliche Familien) interessieren. Dies ist im Zusammenhang mit Matrizen natürlich.

⁴⁴Diese Darstellung existiert und ist eindeutig gemäss Lemma 9.18.

⁴⁵Die Umschreibung am Ende von (10.11) verbindet die Einsteinsche Summenkonvention (10.2) mit der Vereinbarung vor Beispiel 9.2.

⁴⁶Aus diesem Grund scheint es manchen Autoren natürlicher, nicht $\mathcal{S}_{BC}$, sondern $\mathcal{S}_{CB}$ als “Basiswechselmatrix von B nach C ” zu bezeichnen. ymmv

§ 10. BASEN, MATRIZEN, DETERMINANTEN

Beweis von 10.5. (i) Aus

$$v = b_i v^i = c_j w^j = b_i \sigma_j^i w^j \quad (10.16)$$

folgt durch Koeffizientenvergleich (d.h., wegen der Eindeutigkeit der Darstellung von v als Linearkombination von B) unmittelbar Gl. (10.13).

(ii) Es seien τ_i^j für $i, j = 1, \dots, n$ die Einträge von $\mathcal{S}_{CB}$. Dann gilt also

$$b_i = c_j \tau_i^j = b_k \sigma_j^k \tau_i^j \quad (10.17)$$

und daraus folgt durch Koeffizientenvergleich

$$\sigma_j^k \tau_i^j = \delta_i^k := \begin{cases} 1 & \text{falls } k = i \\ 0 & \text{sonst} \end{cases} \quad (10.18)$$

(Kronecker-Delta), gleichbedeutend mit $\mathcal{S}_{BC}\mathcal{S}_{CB} = I_n$. Ebenso folgt $\mathcal{S}_{CB}\mathcal{S}_{BC} = I_n$, und daraus die Behauptung.

(iii) Wir schreiben $\mathcal{S}_{CD} = (\rho_k^j)_{j,k=1,\dots,n}$ und $\mathcal{S}_{BD} = (\tau_k^i)_{i,k=1,\dots,n}$. Dann folgt aus $c_j = b_i \sigma_j^i$ und $d_k = c_j \rho_k^j$ durch Einsetzen $d_k = b_i \sigma_j^i \rho_k^j = b_i \tau_k^i$, wo $\tau_k^i = \sigma_j^i \rho_k^j$ durch Koeffizientenvergleich. Das heisst genau $\mathcal{S}_{BD} = \mathcal{S}_{BC}\mathcal{S}_{CD}$. $\square$

Leichte Varianten:

Lemma 10.6. *Sei V ein endlich-dimensionaler Vektorraum und $B = (b_1, \dots, b_n)$ eine Basis von V .*

(i) *Eine Familie $C = (c_1, \dots, c_n)$ von n Linearkombinationen von B ,*

$$c_j = b_i \sigma_j^i \quad (10.19)$$

ist genau dann eine Basis von V , wenn die Matrix $\mathcal{S} = (\sigma_j^i)_{i,j=1,\dots,n} \in GL(n, K)$, d.h. invertierbar ist.

(ii) *Für jede Matrix $\mathcal{S} \in GL(n, K)$ existiert genau eine Basis C von V so, dass $\mathcal{S}_{BC} = \mathcal{S}$.*

Beweis. Übungsaufgabe $\square$

Problemstellung: Wie berechnen wir die neuen Koordinaten aus den alten, oder äquivalent dazu, wie sieht die alte Basis als Linearkombination der neuen aus?

Idee: Wir nutzen die vom Basisfinden her bekannten Umformungen, angewandt auf die Basiswechselmatrix $\mathcal{S}_{BC}$. Genauer:

1. Die Basis B ist genau dadurch charakterisiert, dass ihre Darstellung in der Basis B sehr einfach ist (nämlich gerade durch die Einheitsmatrix gegeben). Die Aufgabe besteht also darin, die Basis C im Hinblick darauf abzuändern, dass ihre Darstellung in der Basis B möglichst “einfacher” wird.

2. Dabei führen wir (zum Unterschied von S. 58) darüber Buch, welche Umformungen wir vorgenommen haben um “ C zurück zu B zu bringen”.

Beobachtung: Die Operationen II, III, IV entsprechen als *elementare Spaltenumformungen* der Matrixmultiplikation von rechts mit den folgenden “elementaren Spaltenumformungsmatrizen”, zur Abkürzung auch *Elementarmatrizen* genannt (vgl. die Beschreibung der Matrixmultiplikation auf S. 57), Hintereinanderausführung der Matrixmultiplikation. Nämlich:

II. Vertauschen von c_{i_1} mit c_{i_2} :

$$P_{i_1 i_2} := \begin{pmatrix} & i_1 & & i_2 \\ 1 & \cdots & 0 & \cdots & 0 & \cdots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \cdots & 0 & \cdots & 1 & \cdots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \cdots & 1 & \cdots & 0 & \cdots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \cdots & 0 & \cdots & 0 & \cdots & 1 \end{pmatrix} \begin{matrix} \\ i_1 \\ \\ i_2 \\ \\ \end{matrix} \quad (10.20)$$

III. Multiplikation von c_i mit $\lambda \neq 0$:

$$Q_i(\lambda) := \begin{pmatrix} & i \\ 1 & \cdots & 0 & \cdots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \cdots & \lambda & \cdots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \cdots & 0 & \cdots & 1 \end{pmatrix} i \quad (10.21)$$

IV. Addition von λc_{i_1} zu c_{i_2} :

$$R_{i_1 i_2}(\lambda) := \begin{pmatrix} & i_1 & & i_2 \\ 1 & \cdots & 0 & \cdots & 0 & \cdots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \cdots & 1 & \cdots & \lambda & \cdots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \cdots & 0 & \cdots & 1 & \cdots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \cdots & 0 & \cdots & 0 & \cdots & 1 \end{pmatrix} \begin{matrix} \\ i_1 \\ \\ i_2 \\ \\ \end{matrix} \quad (10.22)$$

Jede Elementarmatrix kann als eigenständige Basiswechselmatrix interpretiert werden. Dass sich dabei die Koordinaten jeweils mit der Inversen transformieren, ist bereits in (10.7) und (10.8) sichtbar. Es gilt doch $(P_{i_1 i_2})^{-1} = P_{i_1 i_2}$, $(Q_i(\lambda))^{-1} = Q_i(1/\lambda)$ und $(R_{i_1 i_2}(\lambda))^{-1} = R_{i_1 i_2}(-\lambda)$.

Lösung: Stelle mit Hilfe des auf S. 58 beschriebenen Gaussverfahrens Stufenform her. Dabei müssen alle Stufen Höhe 1 haben, sonst wäre C keine Basis gewesen, genauer: Der von C aufgespannte Unterraum hätte Dimension $\not\leq n$. Dies folgt aus den gleichen Überlegungen wie unter Gl. (10.9), mit dem Unterschied, dass wir nicht bezüglich der Standardbasis des K^n , sondern der vorgegebenen Basis B von V argumentieren. Wir können dann durch Umformungen vom Typ IV auch unterhalb der Diagonalen alles freiräumen und (entweder davor oder danach) die Pivot-Elemente mit Umformungen vom Typ II auf 1 normieren.

Ergebnis: Wir haben $\mathcal{S}_{BC}$ auf die Einheitsmatrix umgeformt. Das Produkt der benutzten Spaltenumformungen ist gerade die inverse Matrix $\mathcal{S}_{CB}$.

Beispiel: Die Buchführung besteht in der gleichzeitigen Anwendung der Spaltenumformungen auf die darunter geschriebene Einheitsmatrix. Dass die angegebene Matrix tatsächlich einen Basiswechsel darstellt, mithin invertierbar ist, “entdecken” wir im Laufe der Rechnung. (Determinanten geben ein a priori Kriterium, siehe unten.)⁴⁷

$$\begin{array}{ccc}
 \begin{pmatrix} 1 & 1 & 2 \\ 0 & 1 & 2 \\ 1 & 1 & -1 \\ \hline 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} & \xrightarrow{c_2 \leftarrow c_2 - c_1} & \begin{pmatrix} 1 & 0 & 2 \\ 0 & 1 & 2 \\ 1 & 0 & -1 \\ \hline 1 & -1 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} & \xrightarrow{c_3 \leftarrow c_3 - 2c_1} & \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 2 \\ 1 & 0 & -3 \\ \hline 1 & -1 & -2 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} & \xrightarrow{c_3 \leftarrow c_3 - 2c_2} \\
 \\
 \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 1 & 0 & -3 \\ \hline 1 & -1 & 0 \\ 0 & 1 & -2 \\ 0 & 0 & 1 \end{pmatrix} & \xrightarrow{c_3 \leftarrow -c_3/3} & \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 1 & 0 & 1 \\ \hline 1 & -1 & 0 \\ 0 & 1 & \frac{2}{3} \\ 0 & 0 & -\frac{1}{3} \end{pmatrix} & \xrightarrow{c_1 \leftarrow c_1 - c_3} & \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \\ \hline 1 & -1 & 0 \\ -\frac{2}{3} & 1 & \frac{2}{3} \\ \frac{1}{3} & 0 & -\frac{1}{3} \end{pmatrix}
 \end{array} \tag{10.23}$$

Variante: Der obige Algorithmus berechnet $\mathcal{S}_{CB}$ als “Rechtsinverse” von $\mathcal{S}_{BC}$ durch Spaltenumformungen entsprechenden Änderungen der “neuen” Basis, C . Stattdessen könnte man auch auf die Idee kommen, die “alte” Basis, B , abzuändern, um die Darstellung von C zu vereinfachen. Hierzu stellen wir fest:

1. Die Basis C ist genau dadurch charakterisiert, dass in ihr die Darstellung der Basis C sehr einfach ist (nämlich durch die Einheitsmatrix gegeben).
2. Transformationen von B vom Typ II, III, IV auf S. 58 entsprechen auf der Ebene der Basiswechselmatrix *elementaren Zeilenumformungen*, und diese wiederum der Multiplikation mit den Elementarmatrizen (10.20), (10.21) (10.22), *von links*.
3. Wenn wir bei der Herstellung der *Zeilenstufenform* einer Nullzeile begegneten, so wäre dies ein Signal dafür, dass die entsprechende Linearkombination von B sich nicht durch C darstellen liesse, dass also C keine Basis gewesen wäre.
4. Wegen der Gleichheit von Links- und Rechtsinverser (unter der Voraussetzung, dass beide existieren) muss das Produkt der elementaren Zeilenumformungen wieder $\mathcal{S}_{CB}$ ergeben.

⁴⁷Erklärung der Notation: Im ersten Schritt wird die erste Spalte von der zweiten abgezogen. Dies entspricht der Multiplikation mit $R_{12}(-1)$ (von rechts), und dem Basiswechsel $(\tilde{c}_1, \tilde{c}_2, \tilde{c}_3) = (c_1, c_2 - c_1, c_3)$. Das heisst nach diesem Schritt geben die Spalten der oberen Matrix die Darstellung der Basis $(\tilde{c}_1, \tilde{c}_2, \tilde{c}_3)$ bezüglich der Basis (b_1, b_2, b_3) .

Im Beispiel:⁴⁸

$$\begin{aligned}
 & \left(\begin{array}{ccc|ccc} 1 & 1 & 2 & 1 & 0 & 0 \\ 0 & 1 & 2 & 0 & 1 & 0 \\ 1 & 1 & -1 & 0 & 0 & 1 \end{array} \right) \xrightarrow{b_1 \rightarrow b_1 - b_3} \left(\begin{array}{ccc|ccc} 1 & 1 & 2 & 1 & 0 & 0 \\ 0 & 1 & 2 & 0 & 1 & 0 \\ 0 & 0 & -3 & -1 & 0 & 1 \end{array} \right) \xrightarrow{b_2 \rightarrow b_2 - b_1} \\
 & \left(\begin{array}{ccc|ccc} 1 & 0 & 0 & 1 & -1 & 0 \\ 0 & 1 & 2 & 0 & 1 & 0 \\ 0 & 0 & -3 & -1 & 0 & 1 \end{array} \right) \xrightarrow{b_3 \rightarrow b_3 + 2b_2/3} \left(\begin{array}{ccc|ccc} 1 & 0 & 0 & 1 & -1 & 0 \\ 0 & 1 & 0 & -\frac{2}{3} & 1 & \frac{2}{3} \\ 0 & 0 & -3 & -1 & 0 & 1 \end{array} \right) \xrightarrow{b_3 \rightarrow -b_3/3} \\
 & \xrightarrow{\quad} \left(\begin{array}{ccc|ccc} 1 & 0 & 0 & 1 & -1 & 0 \\ 0 & 1 & 0 & -\frac{2}{3} & 1 & \frac{2}{3} \\ 0 & 0 & 1 & \frac{1}{3} & 0 & -\frac{1}{3} \end{array} \right)
 \end{aligned} \tag{10.24}$$

Probe:

$$\begin{pmatrix} 1 & 1 & 2 \\ 0 & 1 & 2 \\ 1 & 1 & -1 \end{pmatrix} \begin{pmatrix} 1 & -1 & 0 \\ -\frac{2}{3} & 1 & \frac{2}{3} \\ \frac{1}{3} & 0 & -\frac{1}{3} \end{pmatrix} = \begin{pmatrix} 1 & -1 & 0 \\ -\frac{2}{3} & 1 & \frac{2}{3} \\ \frac{1}{3} & 0 & -\frac{1}{3} \end{pmatrix} \begin{pmatrix} 1 & 1 & 2 \\ 0 & 1 & 2 \\ 1 & 1 & -1 \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \tag{10.25}$$

Determinantenkalkül

Die Anwendung des Basiswechselalgorithmus auf beliebige Matrizen $\mathcal{S} \in \text{Mat}_n(K)$ schliesst “erfolgreich” (d.h., mit der inversen Matrix $\mathcal{S}^{-1}$), falls $\mathcal{S}$ ursprünglich invertierbar war. Andernfalls scheitert er, wenn nämlich die Stufenform eine Nullspalte enthält (oder die Zeilenstufenform eine Nullzeile).⁴⁹ In der Praxis ist dies schon recht effizient, für theoretische Zwecke ist ein a priori, notwendiges und hinreichendes, Kriterium für Invertierbarkeit und eine explizite Formel für die Inverse manchmal ebenso nützlich.

Vorüberlegungen: Wir formulieren zunächst die obigen Ergebnisse um zu Aussagen über $\text{Mat}_n(K)$. Die Interpretation von $GL(n, K)$ als Basiswechselmatrizen eines n -dimensionalen Vektorraum ist dabei nicht mehr essentiell, aber immer noch sehr nützlich, und man darf insbesondere an $V = K^n$ denken.

Lemma 10.7 (Äquivalenz von Links- und Rechtsinvertierbarkeit⁵⁰). *Angenommen, die Matrizen $\mathcal{S} = (\sigma_j^i)_{i,j=1,\dots,n} \in \text{Mat}_n(K)$ und $\mathcal{T} = (\tau_j^i)_{i,j=1,\dots,n} \in \text{Mat}_n(K)$ erfüllen $\mathcal{S}\mathcal{T} = I_n$. Dann gilt auch $\mathcal{T}\mathcal{S} = I_n$.*

Beweis. Es sei $(e_i)_{i=1,\dots,n}$ die Standardbasis des K^n aus Gl. (9.25) und $C = (c_j)_{j=1,\dots,n}$ die Familie von Vektoren, definiert durch⁵¹

$$c_j = e_i \sigma_j^i \tag{10.26}$$

⁴⁸Erklärung der Notation: Im ersten Schritt wird die erste Zeile von der dritten abgezogen. Dies entspricht der Multiplikation mit $R_{31}(-1)$ (von links), und dem Basiswechsel $(b_1, b_2, b_3) = (\tilde{b}_1 - \tilde{b}_3, \tilde{b}_2, \tilde{b}_3)$. Das heisst nach diesem Schritt geben die Spalten der linken Matrix die Darstellung der Basis (c_1, c_2, c_3) bezüglich der Basis $(\tilde{b}_1, \tilde{b}_2, \tilde{b}_3)$.

⁴⁹Dies folgt speziell aus Lemma 10.6: Ist die Matrix nicht invertierbar, so ist C keine Basis, und daher weder ein Erzeugendensystem ($\Rightarrow$ Nullzeile in der Zeilenstufenform) noch linear unabhängig ($\Rightarrow$ Nullspalte in der Stufenform).

⁵⁰Achtung, gilt so nur für (endliche!) Matrizen.

⁵¹Achtung, psychologischer Trick: c_j ist einfach die j -te Spalte von $\mathcal{S}$.

§ 10. BASEN, MATRIZEN, DETERMINANTEN

Dann gilt wegen $\sigma_j^i \tau_k^j = \delta_k^i$, dass

$$e_i = c_j \tau_i^j \quad (10.27)$$

Daraus folgt, dass C ein Erzeugendensystem von K^n ist, und damit aus Korollar (9.25), dass C eine Basis ist. Wegen Lemma 10.6 ist $\mathcal{S}$ invertierbar. Dann folgt (Standardüberlegung) aus $\mathcal{ST} = I_n$, dass $\mathcal{T} = \mathcal{S}^{-1}$ und daraus $\mathcal{TS} = I_n$. $\square$

Lemma 10.8 (Kriterium für Invertierbarkeit). *Eine Matrix ist genau dann invertierbar, wenn sie Produkt von (endlich vielen) Elementarmatrizen der Form (10.20), (10.21) und (10.22) ist.*

Beweis. Elementarmatrizen sind invertierbar (s. S. 62), und daher auch ihre Produkte (s. Bemerkung unter Beispiel 5.4). Umgekehrt stellt der Basiswechselalgorithmus die inverse Matrix einer invertierbaren Matrix als Produkt von Elementarmatrizen dar, so dass die ursprüngliche Matrix das Produkt der inversen Elementarmatrizen (in der umgekehrten Reihenfolge) ist. Beachte: Die Darstellung als Produkt von Elementarmatrizen ist nicht eindeutig. Die Elementarmatrizen können hier entweder als Zeilen- oder Spaltenumformungen aufgefasst werden. $\square$

Lemma 10.9 (Kriterium für Nicht-Invertierbarkeit). *Eine (quadratische) Matrix ist genau dann nicht invertierbar, wenn sie Produkt einer (quadratischen) Matrix in Stufenform mit mindestens einer Nullspalte mit Elementarmatrizen “von rechts” ist.*

Beweis. Die Multiplikation einer Matrix mit Elementarmatrizen von rechts ändert nicht den von den Spalten aufgespannten Unterraum (dies war genau die charakteristische Eigenschaft der Umformungen auf S. 58). Ausgehend von einer Matrix in Stufenform mit Nullspalte(n) kann so keine invertierbare Matrix entstehen. Umgekehrt bringt der Gaußsche Algorithmus eine nicht-invertierbare Matrix auf Stufenform mit Nullspalte(n) durch Multiplikation mit Elementarmatrizen (von rechts). Durch Inversion der Elementarmatrizen folgt die Behauptung. $\square$

Lemma 10.10 ((Nicht-)Invertierbarkeit und Multiplikation von Matrizen). *Für alle $\mathcal{S}, \mathcal{T} \in \text{Mat}_n(K)$ gilt: $\mathcal{ST}$ ist genau dann invertierbar, wenn sowohl $\mathcal{S}$ als auch $\mathcal{T}$ invertierbar sind.*

Beweis. Die direkte Implikation ($\mathcal{S}, \mathcal{T} \in GL(n, K) \Rightarrow \mathcal{ST} \in GL(n, K)$) folgt aus der Gruppenstruktur von $GL(n, K)$ (Abgeschlossenheit unter Multiplikation). Umkehrung: Ist $\mathcal{T}$ nicht invertierbar, so ist nach Kriterium 10.9 $\mathcal{T}$ Produkt einer Matrix in Stufenform mit Nullspalte(n) mit Elementarmatrizen. Dann lässt sich $\mathcal{ST}$ durch Spaltenumformungen in eine Matrix in Stufenform mit Nullspalte(n) bringen und ist damit nicht invertierbar. Analog geht's (mit Zeilenstufenform), falls $\mathcal{S}$ nicht invertierbar ist. $\square$

Positionierung: Die Determinante ist (wie z.B. auch das Kreuzprodukt Gl. (8.19)) ein Beispiel für eine “alternierende Multilinearform”. Diese gehören formal in den Kontext der multi-linearen Algebra und spielen eine wichtige Rolle in der Integrationstheorie und der Vielteilchenquantenmechanik. Wir definieren sie ad hoc und

stellen den Zusammenhang zur Invertierbarkeit von Matrizen her über die Invarianz unter den elementaren Spaltenumformungen. Wir identifizieren hierfür eine Matrix

$$\mathcal{S} = (\sigma_j^i)_{i,j=1,\dots,n} \text{ als Tupel } \mathcal{S} = (s_1, \dots, s_n) \text{ ihrer Spaltenvektoren } s_j = \begin{pmatrix} \sigma_j^1 \\ \vdots \\ \sigma_j^n \end{pmatrix} \in K^n.$$

Definition 10.11. Die *Determinante* ist die Abbildung

$$\det : \text{Mat}_n(K) \rightarrow K \quad (10.28)$$

welche rekursiv in n wie folgt definiert ist:⁵²

- Für $n = 1$: $\det(\sigma_1^1) := \sigma_1^1 \in K$
- Für $n > 1$:

$$\det \mathcal{S} = \det(s_1, \dots, s_n) := \sum_{j=1}^n (-1)^{1+j} \sigma_j^1 \det \mathcal{S}_1^j \quad (10.29)$$

wobei $\mathcal{S}_1^j$ die Matrix ist, die aus $\mathcal{S}$ durch Entfernen der ersten Zeile und j -ten Spalte entsteht, anschaulich:

$$\mathcal{S}_1^j = \begin{pmatrix} \sigma_1^1 & \dots & \sigma_j^1 & \dots & \sigma_n^1 \\ \sigma_1^2 & \dots & \sigma_j^2 & \dots & \sigma_n^2 \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ \sigma_1^n & \dots & \sigma_j^n & \dots & \sigma_n^n \end{pmatrix} \quad (10.30)$$

und $\det \mathcal{S}_1^j$ durch Rekursionsansatz bereits definiert ist, da $\mathcal{S}_1^j \in \text{Mat}_{n-1}(K)$.

Beispiel.

$$\begin{aligned} \det \begin{pmatrix} a & b \\ c & d \end{pmatrix} &= ad - bc \\ \det \begin{pmatrix} \sigma_1^1 & \sigma_2^1 & \sigma_3^1 \\ \sigma_1^2 & \sigma_2^2 & \sigma_3^2 \\ \sigma_1^3 & \sigma_2^3 & \sigma_3^3 \end{pmatrix} &= \sigma_1^1 \det \begin{pmatrix} \sigma_2^2 & \sigma_3^2 \\ \sigma_2^3 & \sigma_3^3 \end{pmatrix} - \sigma_2^1 \det \begin{pmatrix} \sigma_1^2 & \sigma_3^2 \\ \sigma_1^3 & \sigma_3^3 \end{pmatrix} + \sigma_3^1 \det \begin{pmatrix} \sigma_1^2 & \sigma_2^2 \\ \sigma_1^3 & \sigma_2^3 \end{pmatrix} \\ &= \sigma_1^1 \sigma_2^2 \sigma_3^3 - \sigma_1^1 \sigma_3^2 \sigma_2^3 - \sigma_2^1 \sigma_1^2 \sigma_3^3 + \sigma_2^1 \sigma_3^2 \sigma_1^3 + \sigma_3^1 \sigma_1^2 \sigma_2^3 - \sigma_3^1 \sigma_2^2 \sigma_1^3 \end{aligned} \quad (10.31)$$

(vgl. Spatprodukt (8.22))

Theorem 10.12. $\mathcal{S} \in \text{Mat}_n(K)$ ist genau dann invertierbar, wenn $\det \mathcal{S} \neq 0$.

Eigenschaften der Determinante, die für den Beweis von Theorem 10.12 ausreichen.

Lemma 10.13. (i) Die Determinante ist linear in den Spalten: Für alle $i = 1, \dots, n$ gilt, wenn $s_i = t_i + \alpha u_i$ mit $t_i, u_i \in K^n$ und $\alpha \in K$,

$$\det(s_1, \dots, t_i + \alpha u_i, \dots, s_n) = \det(s_1, \dots, t_i, \dots, s_n) + \alpha \det(s_1, \dots, u_i, \dots, s_n) \quad (10.32)$$

⁵²Für $n = 0$: $\det(\emptyset) = 1$.

§ 10. BASEN, MATRIZEN, DETERMINANTEN

(ii) Die Determinante ist total anti-symmetrisch und alternierend bezüglich Spalten:
Für alle $1 \leq i_1 < i_2 \leq n$ gilt

$$\det(\dots, s_{i_1}, \dots, s_{i_2}, \dots) = -\det(\dots, s_{i_2}, \dots, s_{i_1}, \dots) \quad (\text{anti-symmetrisch}) \quad (10.33)$$

$$= 0 \quad \text{falls } s_{i_1} = s_{i_2} \quad (\text{alternierend}) \quad (10.34)$$

(iii) $\det I_n = 1$

Beweis. (i) Bei der rekursiven Auswertung der Formel (10.29) ist der Summand für $j = i$, nämlich $(-1)^{1+i} \sigma_i^1 \det \mathcal{S}_1^i$, offenbar linear in s_i , da $\mathcal{S}_1^i$ nicht mehr von s_i abhängt. Die Summanden für $j \neq i$ hängen von s_i nur über $\det \mathcal{S}_1^j$ ab, und diese Abhängigkeit ist nach Rekursionsvoraussetzung linear.

(ii) Wir bemerken zunächst, dass (10.33) und (10.34) (fast) äquivalent sind: Mit Hilfe der Linearität folgt unter der Gültigkeit von (10.34)

$$\begin{aligned} 0 &= \det(\dots, s_{i_1} + s_{i_2}, \dots, s_{i_1} + s_{i_2}, \dots) = \underbrace{\det(\dots, s_{i_1}, \dots, s_{i_1}, \dots)}_{=0} + \\ &\quad \det(\dots, s_{i_1}, \dots, s_{i_2}, \dots) + \det(\dots, s_{i_2}, \dots, s_{i_1}, \dots) + \underbrace{\det(\dots, s_{i_2}, \dots, s_{i_2}, \dots)}_{=0} \end{aligned} \quad (10.35)$$

d.h. $\det(\dots, s_{i_1}, \dots, s_{i_2}, \dots) = -\det(\dots, s_{i_2}, \dots, s_{i_1}, \dots)$. Umgekehrt folgt (10.34) aus (10.33) via $\det(\dots, s_{i_1}, \dots, s_{i_1}, \dots) = -\det(\dots, s_{i_1}, \dots, s_{i_1}, \dots) = 0$, wenn auch nur für $2 \neq 0 \in K$.

Nun folgt (10.34) für benachbarte Spalten ($i_2 = i_1 + 1$) wieder rekursiv aus (10.29): Ist $s_{i_1} = s_{i_1+1}$, so gilt $\mathcal{S}_1^{i_1} = \mathcal{S}_1^{i_1+1}$ und die Summanden für $j = i_1$ und $i_1 + 1$ heben sich gegenseitig auf, während die übrigen Summanden nach Rekursionsvoraussetzung verschwinden, da sie zwei gleiche Spalten enthalten. Also gilt auch (10.33) für $i_2 = i_1 + 1$. Weiter entfernte gleiche Spalten können dann durch Tauschen mit dazwischenliegenden bis auf Vorzeichen nebeneinander gebracht werden. Dies impliziert (10.34) allgemein.

(iii) Dies folgt unmittelbar rekursiv aus der Definition. $\square$

Lemma 10.14. (i) Bei Multiplikation mit Elementarmatrizen (10.20), (10.21), (10.22) gilt:

$$\begin{aligned} \det(\mathcal{S}P_{i_1 i_2}) &= -\det \mathcal{S} \\ \det(\mathcal{S}Q_i(\lambda)) &= \lambda \det \mathcal{S} \\ \det(\mathcal{S}R_{i_1 i_2}(\lambda)) &= \det \mathcal{S} \end{aligned} \quad (10.36)$$

(ii) $\det(P_{i_1 i_2}) = -1$, $\det(Q_i(\lambda)) = \lambda$, $\det(R_{i_1 i_2}(\lambda)) = 1$.

(iii) $\det(\mathcal{S}\mathcal{E}) = \det \mathcal{S} \cdot \det \mathcal{E}$ für alle $\mathcal{S} \in \text{Mat}_n(K)$ und alle Elementarmatrizen $\mathcal{E}$.

(iv) Für eine Matrix in unterer Dreiecksform,

$$\mathcal{S} = \begin{pmatrix} \lambda_1 & 0 & \cdots & 0 \\ * & \lambda_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ * & * & \cdots & \lambda_n \end{pmatrix} \quad (10.37)$$

gilt $\det \mathcal{S} = \lambda_1 \cdot \lambda_2 \cdots \lambda_n$

Beweis. (i) $P_{i_1 i_2}$ vertauscht zwei Spalten $\rightsquigarrow$ nutze Antisymmetrie. $Q_i(\lambda)$ multipliziert eine Spalte mit $\lambda \in K^\times \rightsquigarrow$ nutze Linearität. Für $R_{i_1 i_2}(\lambda)$ (Addition eines Vielfaches einer Spalte zu einer anderen) nutzen wir Linearität und Alternierung:

$$\det(\dots, s_{i_1}, \dots, s_{i_2} + \lambda s_{i_1}, \dots) = \det(\dots, s_{i_1}, \dots, s_{i_2}, \dots) + \underbrace{\lambda \det(\dots, s_{i_1}, \dots, s_{i_1}, \dots)}_{=0} \quad (10.38)$$

(ii) Folgt mit $\mathcal{S} = I_n$ aus (i) und $\det I_n = 1$.

(iii) Unmittelbar aus (i) und (ii).

(iv) Unmittelbar rekursiv aus der Definition. $\square$

Beweis von Theorem 10.12. Ist $\mathcal{S}$ invertierbar, so ist $\mathcal{S}$ nach Lemma 10.8 Produkt von endlich vielen Elementarmatrizen. Nach Lemma 10.14 (iii) (mehrmals angewandt) ist $\det \mathcal{S}$ dann das Produkt der Determinanten der Elementarmatrizen, die nach 10.14 (ii) alle nicht null sind.

Ist umgekehrt $\mathcal{S}$ nicht invertierbar, so ist $\det \mathcal{S}$ nach Lemma 10.9 und Lemma 10.14 bis auf das Produkt der Elementardeterminanten gleich der Determinante einer Matrix in Stufenform mit Nullspalte(n). Diese verschwindet nach 10.14 (iv), bzw. alternativ wegen der Linearität. $\square$

Konsequenzen: Wie eingangs angekündigt ist die wesentliche Eigenschaft der Determinante ihr Transformationsverhalten unter Spaltenumformungen. Insbesondere ist sie durch ihre Werte auf den Elementarmatrizen bereits vollständig bestimmt.

Proposition 10.15 (Eindeutigkeit). *Sei $\Delta : \text{Mat}_n(K) \rightarrow K$ eine Abbildung, welche die Eigenschaften (i), (ii) und (iii) aus Lemma 10.13 erfüllt. Dann gilt $\Delta = \det$.*

Beweis. Der Beweis von Lemma 10.14 (und Theorem 10.12) lässt sich völlig ohne direkten Verweis auf (10.29) führen (inkl. Eigenschaft (iv), als Übungsaufgabe). Δ und $\det$ haben also die gleichen Werte auf allen Elementarmatrizen, und damit auf ganz $\text{Mat}_n(K)$. $\square$

Bemerkung. Die Eigenschaften (i), (ii), and (iii) aus Lemma 10.13 reichen zur *Definition* der Determinante nicht aus, da die Darstellung einer Matrix als Produkt von Elementarmatrizen nicht eindeutig ist. Die Unabhängigkeit von der Wahl, und damit die *Existenz* einer Abbildung mit den geforderten Eigenschaften, folgt erst aus einer expliziten Formel.

Proposition 10.16 (Multiplikativität). *Für alle $\mathcal{S}_1, \dots, \mathcal{S}_m \in \text{Mat}_n(K)$ gilt*

$$\det(\mathcal{S}_1 \cdots \mathcal{S}_m) = \det \mathcal{S}_1 \cdots \det \mathcal{S}_m \quad (10.39)$$

Für alle $\mathcal{S} \in GL(n, K)$ gilt

$$\det(\mathcal{S}^{-1}) = (\det \mathcal{S})^{-1} \quad (10.40)$$

§ 10. BASEN, MATRIZEN, DETERMINANTEN

Beweis. Es genügt (10.39) für $m = 2$ zu beweisen. Ist $\det \mathcal{S}_2 = 0$, dann ist $\mathcal{S}_2$ nicht invertierbar, und dann nach Lemma 10.10 auch $\mathcal{S}_1 \mathcal{S}_2$, mithin $\det(\mathcal{S}_1 \mathcal{S}_2) = 0$. Ist $\det \mathcal{S}_2 \neq 0$, so ist $\mathcal{S}_2$ Produkt von Elementarmatrizen und dann folgt die Aussage aus Lemma 10.14 (iii). Gl. (10.40) folgt dann mittels $1 = \det I_n = \det(\mathcal{S} \mathcal{S}^{-1}) = \det \mathcal{S} \cdot \det(\mathcal{S}^{-1})$. $\square$

Proposition 10.17 (Vollständige Entwicklung nach Leibniz). *Es gilt*

$$\det \mathcal{S} = \sum_{\pi \in S_n} \operatorname{sgn}(\pi) \prod_{j=1}^n \sigma_j^{\pi(j)} \quad (10.41)$$

wobei $S_n = \operatorname{Bij}(\{1, 2, \dots, n\})$ die symmetrische Gruppe aus Def. 5.5 bezeichnet, und $\operatorname{sgn} : S_n \rightarrow \{+1, -1\}$ den sog. Signum-Homomorphismus (Erklärungen im Beweis).

Beweis. Wir schreiben die Spalten von $\mathcal{S}$, $s_j = e_{i_j} \sigma_j^{i_j}$ als Linearkombinationen der Standardbasis (9.25) des K^n . Dann folgt aus der Linearität:

$$\det(s_1, \dots, s_n) = \sum_{i_1=1}^n \sum_{i_2=1}^n \cdots \sum_{i_n=1}^n \sigma_1^{i_1} \cdot \sigma_2^{i_2} \cdots \sigma_n^{i_n} \det(e_{i_1}, e_{i_2}, \dots, e_{i_n}) \quad (10.42)$$

Da die Determinante verschwindet, wenn zwei Spalten gleich sind, tragen nur Summanden bei, bei denen die $i_1, i_2, \dots, i_n$ paarweise verschieden sind, wenn also eine Bijektion $\pi \in S_n$ existiert so dass $\forall j = 1, \dots, n : i_j = \pi(j)$ (und jede solche Permutation kommt genau einmal vor). Die Matrix $\mathcal{E}_\pi = (e_{\pi(1)}, \dots, e_{\pi(n)})$ lässt sich durch paarweises Vertauschen (Transposition) von Spalten auf I_n umformen, d.h. ihre Determinante $\operatorname{sgn}(\pi) := \det \mathcal{E}_\pi \in \{+1, -1\}$ und der Beweis ist erbracht. Interpretation: $\operatorname{sgn}(\pi)$ ist $+1$ wenn sich π als Komposition (Hintereinanderausführung) einer geraden Zahl von Transpositionen darstellen lässt, und -1 wenn ungerade. Die Unabhängigkeit von der Darstellung folgt aus der Wohldefiniertheit der Determinante. Es gilt auch: $\operatorname{sgn}(\pi_1 \circ \pi_2) = \operatorname{sgn}(\pi_1) \cdot \operatorname{sgn}(\pi_2)$ (Übungsaufgabe). Alternativ kann man zeigen (vorausgesetzt man weiss, dass sgn wohldefiniert ist), dass die Formel (10.41) die Eigenschaften aus Lemma 10.13 erfüllt und daher nach Proposition 10.15 gleich $\det$ sein muss. $\square$

Korollar 10.18. $\det \mathcal{S} = \det \mathcal{S}^T$, wobei $\mathcal{S}^T$ die transponierte⁵³ Matrix bezeichnet, die aus $\mathcal{S}$ durch Spiegelung an der Diagonalen entsteht,

$$\mathcal{S} = \begin{pmatrix} \sigma_1^1 & \sigma_2^1 & \cdots & \sigma_n^1 \\ \sigma_1^2 & \sigma_2^2 & \cdots & \sigma_n^2 \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_1^n & \sigma_2^n & \cdots & \sigma_n^n \end{pmatrix} \rightsquigarrow \mathcal{S}^T = \begin{pmatrix} \sigma_1^1 & \sigma_2^1 & \cdots & \sigma_n^1 \\ \sigma_1^2 & \sigma_2^2 & \cdots & \sigma_n^2 \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_1^n & \sigma_2^n & \cdots & \sigma_n^n \end{pmatrix} \quad (10.43)$$

Insbesondere ist die Determinante auch alternierend (und daher anti-symmetrisch) in den Zeilen.

Bemerkung. Beachte, dass die Transposition mit ihrer “Verletzung der Hoch-/Tiefstellungsregeln” etwas aus dem bisherigen Rahmen fällt.

⁵³Sprachmissbrauch!!!

Beweis. In jedem Faktor eines jeden Summanden von Gl. (10.41) gilt: ‘oberer Index’ = π (‘unterer Index’), d.h. äquivalent ‘unterer Index’ = π^{-1} (‘oberer Index’), wobei $\pi^{-1} \in S_n$ die zu π inverse Permutation bezeichnet, und jeder obere Index und jede (inverse) Permutation genau einmal vorkommt. Wegen

$$\operatorname{sgn}(\pi^{-1}) = \det \mathcal{E}_{\pi^{-1}} = \det((\mathcal{E}_{\pi})^{-1}) = (\det \mathcal{E}_{\pi})^{-1} = \det \mathcal{E}_{\pi} = \operatorname{sgn}(\pi) \quad (10.44)$$

folgt

$$\sum_{\pi \in S_n} \operatorname{sgn}(\pi) \prod_{j=1}^n \sigma_j^{\pi(j)} = \sum_{\pi \in S_n} \operatorname{sgn}(\pi) \prod_{i=1}^n \sigma_{\pi^{-1}(i)}^i = \sum_{\xi \in S_n} \operatorname{sgn}(\xi) \prod_{i=1}^n \sigma_{\xi(i)}^i \quad (10.45)$$

wo wir im letzten Schritt noch $\pi = \xi^{-1}$ “substituiert” haben. $\square$

Proposition 10.19 (Allgemeine Entwicklung). *Für jedes $k = 1, \dots, n$ und (alternativ) $l = 1, \dots, n$ gilt*

$$\begin{aligned} \det \mathcal{S} &= \sum_{j=1}^n (-1)^{k+j} \sigma_j^k \det \mathcal{S}_k^j \quad (\text{Entwicklung nach der } k\text{-ten Zeile}) \\ &= \sum_{i=1}^n (-1)^{i+l} \sigma_i^l \det \mathcal{S}_i^l \quad (\text{Entwicklung nach der } l\text{-ten Spalte}) \end{aligned} \quad (10.46)$$

wobei $\mathcal{S}_k^l$ die $(n-1) \times (n-1)$ -Matrix ist, die aus $\mathcal{S}$ durch Entfernen der k -ten Zeile und l -ten Spalte entsteht (vergleiche Bild (10.30)).

Beweis. Folgt aus der Antisymmetrie in Verbindung mit Korollar 10.18. $\square$

Korollar 10.20 (Sonstiges).

- (i) $\det(\mathcal{S}_1 \cdots \mathcal{S}_m) = \det(\mathcal{S}_{\xi(1)} \cdots \mathcal{S}_{\xi(m)})$ für alle $\xi \in S_m$
- (ii) $\det((\mathcal{S}^T)^{-1}) = (\det \mathcal{S})^{-1}$
- (iii) $(\mathcal{S} \cdot \mathcal{T})^T = \mathcal{T}^T \cdot \mathcal{S}^T$.
- (iv) $\mathcal{S}$ invertierbar $\Leftrightarrow \mathcal{S}^T$ invertierbar und $(\mathcal{S}^T)^{-1} = (\mathcal{S}^{-1})^T$ ⁵⁴
- (v) $\det(\alpha \mathcal{S}) = \alpha^n \det \mathcal{S}$ für alle $\alpha \in K$
- (vi) Für eine Matrix $\mathcal{S} \in \operatorname{Mat}_n(K)$ in “unterer Blockform”

$$\mathcal{S} = \begin{pmatrix} \mathcal{S}_1 & 0 & \cdots & 0 \\ * & \mathcal{S}_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ * & * & \cdots & \mathcal{S}_p \end{pmatrix} \quad (10.47)$$

wo $\mathcal{S}_r \in \operatorname{Mat}_{n_r}(K)$ für $r = 1, \dots, p$ und $\sum_{r=1}^p n_r = n$ gilt

$$\det \mathcal{S} = \prod_{r=1}^p \det \mathcal{S}_r \quad (10.48)$$

Analog/bis auf Vorzeichen: “obere Blockform”/“permutierte Blockform”

Beweis. offensichtliche Übungsaufgaben $\square$

⁵⁴ Ich schreibe für diese auch “kontragredient” genannte Matrix gerne $\mathcal{S}^{-T}$. Es gilt $(\mathcal{S} \cdot \mathcal{T})^{-T} = \mathcal{S}^{-T} \cdot \mathcal{T}^{-T}$.

§ 10. BASEN, MATRIZEN, DETERMINANTEN

Anwendung zur expliziten Inversion von Matrizen

Definition 10.21. Die *Adjunkte* einer Matrix $\mathcal{S} = (\sigma_j^i)_{i,j=1,\dots,n} \in \text{Mat}_n(K)$ ist die Matrix $\text{adj}(\mathcal{S}) = (\tilde{\sigma}_j^i)_{i,j=1,\dots,n}$ mit Einträgen

$$\tilde{\sigma}_j^i = (-1)^{i+j} \det \mathcal{S}_j^i \quad (10.49)$$

d.h. bis auf das “schachbrettartige” Vorzeichen steht in der i -ten Zeile und j -ten Spalte von $\text{adj}(\mathcal{S})$ die Determinante der Matrix, die aus $\mathcal{S}$ durch Entfernen der j -ten Zeile und i -ten Spalte entsteht. (Die Indexstellung ist etwas vorteilhafter als bei der Transponierten.)

Theorem 10.22 (Cramersche Regel). *Für alle $\mathcal{S} \in \text{Mat}_n(K)$ gilt*

$$\mathcal{S} \cdot \text{adj}(\mathcal{S}) = \text{adj}(\mathcal{S}) \cdot \mathcal{S} = \det \mathcal{S} \cdot I_n \quad (10.50)$$

Beweis.

$$\sigma_j^i \tilde{\sigma}_k^j = \sum_{j=1}^n (-1)^{k+j} \sigma_j^i \det \mathcal{S}_k^j \quad (10.51)$$

ist gerade die Entwicklung nach der k -ten Zeile der Determinante derjenigen Matrix, die aus $\mathcal{S}$ entsteht, indem man die k -te Zeile mit der i -ten überschreibt (aber die i -te Zeile unangetastet lässt; vgl. Prop. 10.19). Für $i = k$ ist dies gerade die Determinante von $\mathcal{S}$ selbst, sonst kommt eine Zeile doppelt vor und die Determinante verschwindet. Dies ergibt wie behauptet

$$\sigma_j^i \tilde{\sigma}_k^j = \det \mathcal{S} \cdot \delta_k^i = \begin{cases} \det \mathcal{S} & \text{falls } i = k \\ 0 & \text{sonst} \end{cases} \quad (10.52)$$

Der Nachweis von $\tilde{\mathcal{S}}\mathcal{S} = \det \mathcal{S} \cdot I_n$ geht analog. □

Korollar 10.23. *Für alle $\mathcal{S} \in GL(n, K)$ gilt*

$$\mathcal{S}^{-1} = \frac{1}{\det \mathcal{S}} \cdot \text{adj}(\mathcal{S}) \quad (10.53)$$

□

Beispiel. Eine 2×2 Matrix $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$ ist genau dann invertierbar, wenn $ad - bc \neq 0$, und in diesem Fall ist

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix}^{-1} = \frac{1}{ad - bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix} \quad (10.54)$$

Man mache die Probe jeden Morgen zum Frühstück!

Geometrische Deutung: Eine Basis $B = (b_1, \dots, b_n)$ des euklidischen linearen Raums $\mathbb{R}^n$ kann verwendet werden als Abbild eines linearen Koordinatensystems auf dem “physikalischen Raum” aus § 8. In dieser Interpretation gibt die Determinante der aus den Koordinaten der $b_j = (\beta_j^1, \dots, \beta_j^n)^T$ bezüglich der Standardbasis gebildeten Matrix $\mathcal{B} = (\beta_j^i)$ das *orientierte euklidische Volumen* des von den b_j aufgespannten n -dimensionalen Parallelotops (“Hyperspat”), vgl. Abbildung 8.2. Man kann auch sagen: $\det \mathcal{B}$ ist der “Rückzugsfaktor” der euklidischen Volumenform auf das durch B definierte Koordinatensystem.

§ 11 Lineare Abbildungen

Die Aussagen zur Existenz von Basen und die Beschreibung von Basiswechseln beleuchten schon einigermaßen die “innere Struktur” eines festgehaltenen Vektorraums. Zur Strukturtheorie gehört nun weiterhin

1. der Vergleich verschiedener Vektorräume mittels “strukturerhaltender Abbildungen” (Morphismen);
2. die Konstruktion “neuer Vektorräume aus alten”, und die Beschreibung der damit einhergehenden “äusseren” Strukturen, wie etwa der Beziehungen zwischen einem Vektorraum und seinen “Bauelementen”.

Diese Ideen spielten bereits eine Rolle bei der Erweiterung der Rechenbereiche in Kapitel 2, und sind, geeignet interpretiert, auch zentral für den Vergleich mit realen Gegebenheiten, wie er in § 8 angedeutet wurde. Wie zuvor ist jetzt K ein festgehaltener Grundkörper, und man verpasst wenig, wenn man an $K = \mathbb{Q}, \mathbb{R}$ oder $\mathbb{C}$ denkt. Viele Aussagen sind nur für endlich-dimensionale Vektorräume wirklich schlagkräftig; für unendlich-dimensionale Vektorräume benötigt man weitere (topologische) Struktur, siehe § 15 sowie HöMa 2 und 3.

Definition 11.1. Es seien V und W K -Vektorräume. Eine Abbildung $\Phi : V \rightarrow W$ (sc. der zugrundeliegenden Mengen) heisst *linear* falls für alle $v, v_1, v_2 \in V$ und $\alpha \in K$ gilt:

- (i) $\Phi(v_1 + v_2) = \Phi(v_1) + \Phi(v_2)$ (Additivität)
- (ii) $\Phi(\alpha \cdot v) = \alpha \cdot \Phi(v)$ (Homogenität)

Hier stehen links die Operationen mit V und rechts die mit W . Eine lineare Abbildung heisst auch *Homomorphismus* von Vektorräumen. Wir schreiben $\text{Hom}(V, W)$ (zur Betonung manchmal auch $\text{Hom}_K(V, W)$) und reden von K -linearen Abbildungen von K -Vektorräumen) für die Menge der linearen Abbildungen von V nach W .

Übungsaufgabe. Eine Abbildung $\Phi : V \rightarrow W$ ist genau dann linear, wenn für alle $v_1, v_2 \in V$ und $\alpha \in K$ gilt: $\Phi(v_1 + \alpha v_2) = \Phi(v_1) + \alpha \Phi(v_2)$.

Lemma 11.2. Sei $\Phi : V \rightarrow W$ linear. Dann gilt

- (i) $\forall v \in V : \Phi(-v) = -\Phi(v)$
- (ii) $\Phi(0_V) = 0_W$
- (iii) $\forall v_1, \dots, v_k \in V, \alpha_1, \dots, \alpha_k \in K : \Phi\left(\sum_{i=1}^k \alpha_i v_i\right) = \sum_{i=1}^k \alpha_i \Phi(v_i)$

Beweis. (i) folgt am schnellsten mit Hilfe Lemma 9.3 (ii) aus der Homogenität (Def. 11.1 (ii)): $\Phi(-v) = \Phi((-1) \cdot v) = (-1) \cdot \Phi(v) = -\Phi(v)$. (ii) folgt aus (i) in Verbindung mit der Tatsache, dass V nicht leer ist, oder alternativ mit Lemma 9.3 (i) aus $\Phi(0_V) = \Phi(0 \cdot 0_V) = 0 \cdot \Phi(0_V) = 0_W$. (iii) per Induktion und Assoziativität. $\square$

Beispiele 11.3. 1. Für alle Vektorräume V, W ist $\Phi : V \rightarrow W, \Phi(v) := 0_W$ für alle $v \in V$ linear, genannt die Nullabbildung. $\text{Hom}(V, W)$ ist also in jedem Fall (auch für $W = \{0\}$) nicht leer.

2. Die Identitätsabbildung $\text{id}_V : V \rightarrow V$ ist linear.

§ 11. LINEARE ABBILDUNGEN

3. Für jedes $i = 1, \dots, n$ ist $\begin{pmatrix} v^1 \\ \vdots \\ v^n \end{pmatrix} \mapsto v^i$ eine lineare Abbildung $K^n \rightarrow K$.

4. Ist V ein (nicht notwendig endlich-dimensionaler) Vektorraum und B eine Basis von V , so ist für jedes $b \in B$ die Abbildung $V \rightarrow K$,

$$v \mapsto \text{'Koeffizient } v^b \in K \text{ in der Darstellung von } v \text{ als Linearkombination von } B' \quad (11.1)$$

linear. Hierbei wird vereinbart, dass $v^b = 0$ falls b nicht in der Darstellung von v vorkommt (oder es wird darauf bestanden, dass b auf jeden Fall zu verwenden ist). Wegen Lemma 9.18 ist dies wohldefiniert und linear. Beachte: Die Abbildung (11.1) hängt im Allgemeinen nicht nur von b , sondern von der gesamten Basis B ab. Wird $\tilde{b} \in B \setminus \{b\}$ durch $\tilde{b} + \lambda b$ ersetzt, so ändert sich der Koeffizient v^b zu $v^b - \lambda v^{\tilde{b}}$ (siehe Diskussion zum Basiswechsel in § 10).

5. Ist $V = K^n$, $W = K^m$ und $\mathcal{A} = (a_{ij}^j)_{\substack{j=1,\dots,m \\ i=1,\dots,n}} \in \text{Mat}_{m \times n}(K)$ so ist die Abbildung $K^n \rightarrow K^m$,

$$v = \begin{pmatrix} v^1 \\ \vdots \\ v^n \end{pmatrix} \mapsto \mathcal{A} \cdot v = \begin{pmatrix} a_1^1 & \cdots & a_n^1 \\ \vdots & \ddots & \vdots \\ a_1^m & \cdots & a_n^m \end{pmatrix} \begin{pmatrix} v^1 \\ \vdots \\ v^n \end{pmatrix} = \begin{pmatrix} a_1^1 v^1 \\ \vdots \\ a_i^m v^i \end{pmatrix} \quad (11.2)$$

linear. Dies folgt (nach einer geeigneten Umordnung des Alphabets) aus den Rechenregeln für Matrizen aus Proposition 10.2.

6. Auf dem Vektorraum der Funktionen $\mathcal{F}(X, K)$ auf einer beliebigen Menge X ist für jedes $x_* \in X$ die *Auswertungsabbildung* $f \mapsto f(x_*)$ linear. Dies folgt direkt aus der Definition (9.4) der Vektorraumstruktur auf $\mathcal{F}(X, K)$. Für endliche Mengen (oder nach Einschränkung auf $\mathcal{F}^{\text{fin}}(X, K)$) gibt die Auswertungsabbildung bei $x_* \in X$ den Koeffizienten von δ_{x_*} in der Basis $(\delta_x)_{x \in X}$, siehe (9.27), als Beispiel für (11.1).

7. Diese Überlegungen gelten insbesondere für Polynome. Für festes $p_* \in K[x]$ ist zwar die Abbildung $K \rightarrow K$, $x \mapsto p_*(x)$ im Allgemeinen nicht linear. Für festes $x_* \in K$ allerdings ist $K[x] \rightarrow K$, $p \mapsto p(x_*)$ eine lineare Abbildung.

Allgemeine Struktur linearer Abbildungen

- Lineare Abbildungen sind verträglich mit dem Konzept des linearen Unterraums.

Lemma 11.4. *Sei $\Phi \in \text{Hom}(V, W)$ und $U \subset V$ ein UVR. Dann ist auch $\Phi(U) \subset W$ ein UVR. Insbesondere ist das (mengentheoretische) Bild von ganz V (siehe Def. 4.5), ab jetzt geschrieben als*

$$\text{im}(\Phi) = \{w \in W \mid \exists v \in V : \Phi(v) = w\} \subset W \quad (11.3)$$

ein UVR von W .

Beweis. U ist nicht leer, also auch $\Phi(U)$. Ist $w_1 = \Phi(u_1)$ und $w_2 = \Phi(u_2)$, so ist $\Phi(u_1 + \beta u_2) = w_1 + \beta w_2$. $\square$

- Eine neuartige Konstruktion, die von der (additiven) Gruppenstruktur des Wertebereichs abhängt:

Definition 11.5. Sei $\Phi \in \text{Hom}(V, W)$. Dann heit

$$\ker(\Phi) := \{v \in V \mid \Phi(v) = 0_W\} \subset V \quad (11.4)$$

der *Kern* von Φ .

Proposition 11.6. Sei $\Phi \in \text{Hom}(V, W)$. Dann gilt:

- (i) $\ker(\Phi)$ ist ein UVR von V .
- (ii) Φ ist injektiv genau dann wenn $\ker(\Phi) = \{0_V\}$.⁵⁵

Beweis. (i) $\Phi(0_V) = 0_W$, also ist $\ker(\Phi) \neq \emptyset$. Ist $\Phi(v_1) = 0_W$ und $\Phi(v_2) = 0_W$, $\alpha \in K$, so gilt auch $\Phi(v_1 + \alpha v_2) = \Phi(v_1) + \alpha \Phi(v_2) = 0_W$, d.h. $\ker(\Phi)$ ist abgeschlossen unter Addition und Skalarmultiplikation.

(ii) Φ ist per Definition injektiv wenn fr alle $v_1, v_2 \in V$: $\Phi(v_1) = \Phi(v_2) \Rightarrow v_1 = v_2$. Dies ist genau dann der Fall wenn $\Phi(v_1 - v_2) = 0_W \Rightarrow v_1 - v_2 = 0_V$, oder anders gesagt wenn $v_1 - v_2 \in \ker(\Phi) \Rightarrow v_1 - v_2 = 0_V$, d.h. wenn $\ker(\Phi) = \{0_V\}$. $\square$

Definition 11.7. Die Dimension des Bildes heit *Rang* einer linearen Abbildung, die Dimension des Kerns heit *Defekt* (alternativ: Nullitt)

$$\text{rang}(\Phi) := \dim(\text{im}(\Phi)), \quad \text{defect}(\Phi) := \dim(\ker(\Phi)) \quad (11.5)$$

Bemerkung. Wegen 11.4 und 9.26 ist $\dim W < \infty$ eine hinreichende (aber nicht notwendige) Bedingung fr die Endlichkeit von $\text{rang}(\Phi)$. In hnlicher Weise ist wegen 11.6 $\dim V < \infty$ eine hinreichende (aber nicht notwendige) Bedingung fr den Endlichkeit von $\text{defect}(\Phi)$.

- Vektorrume induzieren in natrlicher Weise Strukturen auf den Mengen ihrer Homomorphismen.

Lemma 11.8. (i) Sind V, W K -Vektorrume, so ist $\text{Hom}(V, W)$, ausgerstet mit den punktweise Verknpfungen $(\Phi_1 + \Phi_2)(v) := \Phi_1(v) +_W \Phi_2(v)$, $(\alpha \cdot \Phi)(v) := \alpha \cdot \Phi(v)$, ebenfalls ein K -Vektorraum.

(ii) Ist X ein weiterer K -Vektorraum, $\Phi \in \text{Hom}(V, W)$ und $\Psi \in \text{Hom}(W, X)$, so ist auch $\Psi \circ \Phi : V \rightarrow X$ linear. Es gilt $\Psi \circ (\Phi_1 + \alpha \Phi_2) = \Psi \circ \Phi_1 + \alpha \Psi \circ \Phi_2$, $(\Psi_1 + \beta \Psi_2) \circ \Phi = \Psi_1 \circ \Phi + \beta \Psi_2 \circ \Phi$.

(iii) Ist $\Phi \in \text{Hom}(V, W)$ invertierbar (als Abbildung der zugrundeliegenden Mengen), so ist auch ihre (mengentheoretische) Inverse Φ^{-1} linear.

(iv) $(\text{Hom}(V, V), +, \circ)$ ist ein unitrer Ring (mit Eins gleich Identittsabbildung).

Beweis. (i) und (ii) folgen durch Ausschreiben, (iv) unmittelbar aus (i) und (ii). Die (leicht nicht-triviale) Linearitt der Umkehrabbildung folgt aus der Tatsache, dass $\Phi^{-1}(w_1) + \beta \Phi^{-1}(w_2)$ unter Φ auf $w_1 + \beta w_2$ abgebildet wird, also bereits das mengentheoretische Urbild ist, so dass gilt $\Phi^{-1}(w_1 + \beta w_2) = \Phi^{-1}(w_1) + \beta \Phi^{-1}(w_2)$. $\square$

- Damit assoziierte Sprachregelungen:

⁵⁵Man halte bei dieser dramatischen Vereinfachung wenigstens kurz inne: Fr den Nachweis der Injektivitt einer linearen Abbildung (wie zu gegebener Zeit auch anderer Gruppenhomomorphismen) gengt der Test an einem einzigen Element.

§ 11. LINEARE ABBILDUNGEN

Definition 11.9. Eine lineare Abbildung $\Phi \in \text{Hom}(V, W)$ heisst

- (i) *Epimorphismus*, falls Φ surjektiv ist, d.h. $\text{im}(\Phi) = W$;
- (ii) *Monomorphismus*, falls Φ injektiv ist, d.h. $\ker(\Phi) = \{0_V\}$;
- (iii) *Isomorphismus*, falls Φ bijektiv ist, d.h. $\text{im}(\Phi) = W$ und $\ker(\Phi) = \{0_V\}$;
- (iv) *Endomorphismus*, falls $V = W$. Schreibweise: $\text{End}(V) = \text{Hom}(V, V)$;
- (v) *Automorphismus*, falls $V = W$ und Φ bijektiv. Symbol: $\text{Aut}(V)$ oder auch $GL(V)$, aufgefasst als Einheitengruppe von $(\text{Hom}(V, V), +, \circ)$.

Wir sagen, zwei Vektorräume V und W seien *isomorph* (zueinander), falls ein Isomorphismus $V \rightarrow W$ existiert. Wir schreiben hierfür $V \simeq W$. Beachte, dass es zwischen isomorphen Vektorräumen im Allgemeinen mehrere Isomorphismen gibt, und normalerweise keiner davon ausgezeichnet ist. Einen Isomorphismus, der “nicht von irgendwelchen sonstigen Daten oder Wahlen abhängt” nennt man *natürlich*.⁵⁶

Lemma 11.10. *Isomorphie von Vektorräumen ist eine Äquivalenzrelation.*

Beweis. Wir müssen drei Dinge überprüfen (siehe Def. 3.5). Reflexivität: $\text{id}_V : V \rightarrow V$ ist ein Isomorphismus. Symmetrie: Ist $\Phi : V \rightarrow W$ ein Isomorphismus, so auch $\Phi^{-1} : W \rightarrow V$ (Lemma 11.8 (iii)). Transitivität (Verschärfung von Lemma 11.8 (ii)): Sind $\Phi : V \rightarrow W$ und $\Psi : W \rightarrow X$ Isomorphismen, so auch $\Psi \circ \Phi : V \rightarrow X$. $\square$

• Diese Eigenschaften lassen sich mit Hilfe der in § 9 eingeführten Begriffe bequem(?) umformulieren und testen.

Lemma 11.11. *Sei $\Phi \in \text{Hom}(V, W)$. Dann gilt für jede Menge oder Familie F von Vektoren in V : $\Phi(\text{span}(F)) = \text{span}(\Phi(F))$.*

Bemerkung. Hier bezeichnet für eine Familie $F = (v_i)_{i \in I}$ $\Phi(F) = (\Phi(v_i))_{i \in I}$ die Komposition der indizierenden Abbildung mit Φ .

Beweis. Folgt sofort aus Lemma 11.2 (iii) $\square$

Lemma 11.12. *Für alle $\Phi \in \text{Hom}(V, W)$ sind äquivalent:*

- (i) Φ ist surjektiv.
- (ii) Für jedes Erzeugendensystem E von V ist $\Phi(E)$ ein Erzeugendensystem von W .
- (iii) Es existiert ein Erzeugendensystem E von V so, dass $\Phi(E)$ ein Erzeugendensystem von W ist.

Beweis. Da jeder Vektorraum sich selbst erzeugt, genügt es zu zeigen, dass aus der dritten Aussage die erste, und aus der ersten Aussage die zweite folgt.

(iii) $\Rightarrow$ (i): Ist E ein Erzeugendensystem von V so, dass $\Phi(E)$ ein Erzeugendensystem von W ist, so folgt unter Benutzung von Lemma 11.11

$$\Phi(V) = \Phi(\text{span}(E)) = \text{span}(\Phi(E)) = W \quad (11.6)$$

d.h. Φ ist surjektiv.

⁵⁶ Das ist keine echte Definition, und in Wirklichkeit ist der Begriff um einiges komplizierter. Wir werden ihn nur spärlich verwenden. Im Kontext von Vektorräumen bedeutet Natürlichkeit im Wesentlichen “unabhängig von der Wahl einer Basis”.

(i) $\Rightarrow$ (ii): Ist Φ surjektiv und E ein Erzeugendensystem von V , dann folgt wieder mit Lemma 11.11

$$\text{span}(\Phi(E)) = \Phi(\text{span}(E)) = \Phi(V) = W \quad (11.7)$$

d.h. $\Phi(E)$ ist Erzeugendensystem von W . □

Lemma 11.13. *Für alle $\Phi \in \text{Hom}(V, W)$ sind äquivalent:*

(i) Φ ist injektiv.

(ii) Für jede linear unabhängige Menge $F \subset V$ ist $\Phi(F) \subset W$ linear unabhängig.

Beweis. Ist Φ injektiv, $F \subset V$ linear unabhängig, $w_1, \dots, w_k \in \Phi(F)$ paarweise verschieden, und $\beta_1, \dots, \beta_k \in K$ mit $\sum_i \beta_i w_i = 0_W$, dann existieren $v_1, \dots, v_k \in V$ mit $w_i = \Phi(v_i)$ für alle i und paarweise verschieden. Aus der Linearität folgt $\Phi(\sum_i \beta_i v_i) = 0_W$ und damit aus der Injektivität $\sum_i \beta_i v_i = 0_V$. Wegen der linearen Unabhängigkeit von F muss $\beta_i = 0$ sein für alle $i = 1, \dots, k$. Das heisst $\Phi(F)$ ist linear unabhängig.

Sei umgekehrt für jede linear unabhängige Menge $F \subset V$ auch $\Phi(F) \subset W$ linear unabhängig. Da für alle $v \in V \setminus \{0_V\}$ die ein-elementige Menge $\{v\}$ linear unabhängig ist, kann $\Phi(v)$ nicht gleich 0_W sein. Also ist $\ker(\Phi) = \{0_V\}$, d.h. Φ ist injektiv (s. Proposition 11.6).

Beachte, dass es (natürlich) für die Injektivität nicht ausreicht, wenn das Bild einer (einzelnen) linear unabhängigen Menge linear unabhängig ist. □

Proposition 11.14. *Für alle $\Phi \in \text{Hom}(V, W)$ sind äquivalent:*

(i) Φ ist bijektiv.

(ii) Für jede Basis B von V ist $\Phi(B)$ eine Basis von W .

(iii) Es existiert eine Basis B von V so, dass $\Phi|_B$ injektiv ist und $\Phi(B)$ eine Basis von W .

(iv) Für jede Basis B von V ist $\Phi|_B$ injektiv und $\Phi(B)$ eine Basis von W , d.h. Φ bildet jede Basis von V bijektiv auf eine Basis von W ab.

Beweis. (i) $\Rightarrow$ (ii): Ist Φ bijektiv und B eine Basis von V , so ist $\Phi(B)$ nach Lemma 11.12 Erzeugendensystem von W und nach Lemma 11.13 linear unabhängig, also eine Basis von W .

(ii) $\Rightarrow$ (i): Ist für jede Basis B von V $\Phi(B)$ eine Basis von W , dann ist für eine beliebige Basis B von V $\Phi(B)$ Erzeugendensystem von W , also ist Φ surjektiv. Ist $v \in V \setminus \{0_V\}$, so existiert nach Proposition 9.19 eine Basis B von V mit $v \in B$. Dann kann $\Phi(v) \in \Phi(B)$ nicht gleich 0_W sein, sonst wäre $\Phi(B)$ nicht linear unabhängig. Also ist Φ injektiv.

(i) $\Rightarrow$ (iv): Im Vergleich zu den vorangegangenen Beweisteilen ist nur verschärfend festzuhalten, dass aus der Bijektivität von Φ die Injektivität auf jeder Teilmenge folgt.

(iv) $\Rightarrow$ (iii): Folgt aus der Existenz einer Basis.

§ 11. LINEARE ABBILDUNGEN

(iii) $\Rightarrow$ (i): Ist B eine Basis von V so, dass $\Phi|_B$ injektiv ist und $\Phi(B)$ eine Basis von W , so ist wegen Lemma 11.12 Φ zunächst surjektiv. Ist $\Phi(v) = 0_W$ mit Darstellung $v = \sum_i \alpha_i b_i$ als Linearkombination von B mit b_i paarweise verschieden, so sind wegen der Injektivität von $\Phi|_B$ die $\Phi(b_i)$ paarweise verschieden und daher $\sum_i \alpha_i \Phi(b_i)$ die eindeutige Darstellung von $\Phi(v) = 0_W$ als Linearkombination von $\Phi(B)$. Es folgt $\alpha_i = 0$ für alle i , und daraus $v = 0_V$. Also ist Φ injektiv. $\square$

- Für die Praxis ist auch die Existenz und effiziente Konstruktion von linearen Abbildungen mit Hilfe von Basen von Bedeutung.

Proposition 11.15 (Lineare Fortsetzung). *Es seien V und W zwei K -Vektorräume und $F \subset V$ linear unabhängig. Dann existiert für jede (mengentheoretische) Abbildung $\varphi : F \rightarrow W$ eine lineare Abbildung $\Phi : V \rightarrow W$ so, dass $\Phi|_F = \varphi$. Ist $F = B$ eine Basis von V , so ist diese lineare Fortsetzung eindeutig.*

Beweis. Wir zeigen zunächst die Eindeutigkeit: Sind $\Phi, \tilde{\Phi} \in \text{Hom}(V, W)$ zwei lineare Fortsetzungen von $\varphi : B \rightarrow W$, so gilt für jeden Vektor $v \in V$ mit seiner (eindeutigen) Darstellung $v = \sum_i \alpha_i b_i$ als Linearkombination von B :

$$\Phi(v) = \Phi\left(\sum_i \alpha_i b_i\right) = \sum_i \alpha_i \varphi(b_i) = \tilde{\Phi}\left(\sum_i \alpha_i b_i\right) = \tilde{\Phi}(v) \quad (11.8)$$

Zur Existenz ergänzen wir zunächst (falls nötig) mit Hilfe von Lemma 9.19 F zu einer Basis B von V , und setzen φ von F zu $\tilde{\varphi} : B \rightarrow W$ beliebig fort (z.B. als $\tilde{\varphi}(b) = 0_W$ für $b \in B \setminus F$).⁵⁷ Dann definieren wir für jedes $v \in V$ mit (eindeutiger) Darstellung $v = \sum_i \alpha_i b_i$ als Linearkombination von B :

$$\Phi(v) := \sum_i \alpha_i \tilde{\varphi}(b_i) \quad (11.9)$$

und prüfen (leicht) Wohldefiniertheit und Linearität. Die Eigenschaften $\Phi|_B = \tilde{\varphi}$ und damit $\Phi|_F = \tilde{\varphi}|_F = \varphi$ gelten per Konstruktion. $\square$

Mit anderen Worten: Eine lineare Abbildung $V \rightarrow W$ ist durch ihre Werte auf einer Basis von V bereits vollständig bestimmt.

Korollar 11.16. *Für je zwei Vektorräume V und W sind äquivalent:*

- (i) V und W sind isomorph.
- (ii) Für jede Basis B von V und jede Basis C von W existiert eine Bijektion $B \rightarrow C$.
- (iii) Es existiert eine Basis B von V und eine Basis C von W mit einer Bijektion $B \rightarrow C$.

Insbesondere sind endlich erzeugte Vektorräume dann und nur dann isomorph, wenn $\dim V = \dim W$.

Beweis. Aus Proposition 11.14 folgt, dass jede Basis von V unter einem Isomorphismus auf eine Basis von W abgebildet wird. Dies genügt zum Nachweis von (i) $\Rightarrow$ (iii). Für (i) $\Rightarrow$ (ii) müssen wir noch benutzen, dass je zwei Basen von W gleichmächtig sind, was wir für endlich erzeugte Vektorräume in Theorem 9.20 bewiesen

⁵⁷Siehe Fussnote 14 zur Definition des Begriffs der mengentheoretischen Fortsetzung.

hatten, im Allgemeinen aber wieder von etwas komplizierten mengentheoretischen Axiomen abhängt. Die Implikation (ii) $\Rightarrow$ (iii) ist trivial. Für (iii) $\Rightarrow$ (i) (oder (ii) $\Rightarrow$ (i)) müssen wir noch festhalten, dass die gemäss Proposition 11.15 definierte lineare Fortsetzung einer Bijektion von Basen ein Isomorphismus ist. Dies folgt aber auch sofort aus Proposition 11.14. $\square$

Koordinatendarstellung, Koordinatenwechsel

Theorem 11.17. *Sei V ein endlich-dimensionaler Vektorraum, $n = \dim V \in \mathbb{N}$. Dann ist für jede Basis $B = (b_1, \dots, b_n)$ von V die Abbildung*

$$\Phi_B : K^n \rightarrow V, \quad \begin{pmatrix} v^1 \\ \vdots \\ v^n \end{pmatrix} \mapsto v^i b_i \quad (11.10)$$

ein Isomorphismus von Vektorräumen. Bezeichnet e_i das i -te Element der Standardbasis von K^n (siehe Gl. (9.25)), so gilt $\Phi_B(e_i) = b_i$ für alle $i = 1, \dots, n$. Umgekehrt ist für jeden Isomorphismus $\Phi : K^n \rightarrow V$ $B_\Phi = (\Phi(e_1), \dots, \Phi(e_n))$ eine Basis von V .

Beweis. Umformulierung der Aussagen von Korollar 11.16, ggfs. mit Hilfe von Proposition 11.14, Proposition 11.15, und der Rechnungen aus § 10. $\square$

Definition 11.18. Wir nennen diesen von einer Basis B gestifteten Vektorraumisomorphismus $\Phi_B : K^n \rightarrow V$ die *Parametrisierung* (oder Koordinatisierung) (von V durch K^n mittels B). Die inverse Abbildung $(\Phi_B)^{-1} : V \rightarrow K^n$, heisst *Koordinatenabbildung* (siehe auch Def. 9.23).

Beispiel 11.19. Wie bereits in § 8 angedeutet, spielt die Idee der “Koordinatisierung” und des “Koordinatenwechsels” eine wichtige Rolle in der kinematischen Beschreibung physikalischer Räume. Hierbei erhält man V als “mit (irgendwie ausgezeichneten, üblicherweise euklidischen) Koordinaten $x^1, \dots, x^n$ gegebenen” reellen Raum, und bezeichnet die “Substitution” $x^j = b_i^j v^i$ (mit einer invertierbaren Matrix $\mathcal{B} = (b_i^j)_{i,j=1,\dots,n}$) als “Übergang zu allgemeinen linearen Koordinaten” $v^1, \dots, v^n$. Schreiben wir zur Andeutung der verwendeten Koordinaten $\mathbb{R}_x^n$ für den “parametrisierten”, und $\mathbb{R}_v^n$ für den “parametrisierenden” Vektorraum, so erkennen wir in der Sprache der linearen Algebra $\mathcal{B}$ als Darstellung einer “beliebigen” Basis $(b_1, \dots, b_n)$ als Linearkombination der Standardbasis, $b_i = e_j b_i^j$, also gemäss Theorem 11.17 als lineare Abbildung (s. Gl. 11.2):⁵⁸

$$\begin{aligned} \Phi_B = \mathcal{B} &= \begin{pmatrix} b_1^1 & \cdots & b_n^1 \\ \vdots & \ddots & \vdots \\ b_1^n & \cdots & b_n^n \end{pmatrix} : \mathbb{R}_v^n \rightarrow \mathbb{R}_x^n \\ v = \begin{pmatrix} v^1 \\ \vdots \\ v^n \end{pmatrix} &\mapsto v^i b_i = e_j b_i^j v^i = \mathcal{B} \cdot v = e_j x^j = \begin{pmatrix} x^1 = b_i^1 v^i \\ \vdots \\ x^n = b_i^n v^i \end{pmatrix} \end{aligned} \quad (11.11)$$

⁵⁸Im Kontext der Differentialrechnung verwenden wir zur Parameterisierung noch allgemeinere “krummlinige” Koordinaten, $\xi \in \mathbb{R}_\xi^n$, die mittels differenzierbarer Funktionen $\xi \mapsto x = F(\xi)$ auf den $\mathbb{R}_x^n$ abgebildet werden. Beispiel: Kugelkoordinaten: $x = \rho \sin \vartheta \cos \varphi$, $y = \rho \sin \vartheta \sin \varphi$, $z = \rho \cos \vartheta$.

§ 11. LINEARE ABBILDUNGEN

Ob man nun $v \in V$ als “abstrakten Vektor” und seine Koordinaten $v^i \in K$ als “konkretes Zahlentupel” wertet oder umgekehrt v als “reale Gegebenheit” und die v^i als “imaginäres Hilfsmittel”, ist letztlich eine Gemütsfrage. Wichtig ist in jedem Fall die gedankliche Trennung der beiden Konzepte.

Für lineare Abbildungen zwischen “echt verschiedenen Vektorräumen” lautet die zentrale Aussage wie folgt.

Theorem 11.20. *Sind V und W endlich-dimensional, so ist auch $\text{Hom}(V, W)$ endlich-dimensional und es gilt: $\dim(\text{Hom}(V, W)) = \dim V \cdot \dim W$.*

Beweis. Die Vektorraumstruktur auf $\text{Hom}(V, W)$ wurde in Lemma 11.8 erklärt. Gemäss Lemma 11.15 ist für jede Basis $B = (b_1, \dots, b_n)$ von V die Abbildung $\text{Hom}(V, W) \rightarrow W^{\times n}$, $\Phi \mapsto (\Phi(b_1), \dots, \Phi(b_n))$ eine (zunächst mengentheoretische) Bijektion. Andererseits stiftet jede Basis $C = (c_1, \dots, c_m)$ von W eine Bijektion $W^n \rightarrow \text{Mat}_{m \times n}(K)$, $(w_1, \dots, w_n) \mapsto (w_i^j)_{\substack{j=1, \dots, m \\ i=1, \dots, n}}$. Man prüft leicht, dass die Komposition

$$\begin{aligned} \mathcal{M}_{CB} : \text{Hom}(V, W) &\rightarrow \text{Mat}_{m \times n}(K) \\ \mathcal{M}_{CB}(\Phi) &= \begin{pmatrix} \Phi(b_1)^1 & \cdots & \Phi(b_n)^1 \\ \vdots & \ddots & \vdots \\ \Phi(b_1)^m & \cdots & \Phi(b_n)^m \end{pmatrix} \end{aligned} \quad (11.12)$$

eine lineare Abbildung ist, also (siehe Lemma 11.8 (iii)) ein Vektorraumisomorphismus. Dies impliziert unmittelbar die Aussage. $\square$

Die wesentliche Einsicht des Beweises ist, dass durch Basen B und C der Raum der linearen Abbildungen $\text{Hom}(V, W)$ identifiziert wird mit dem Raum der $m \times n$ -Matrizen mit Einträgen in K . Dies gilt prinzipiell sogar für unendlich-dimensionale Vektorräume (mit geeigneter Interpretation des Raumes der Matrizen). Allerdings sind die Verallgemeinerungen der “Einheitsmatrizen” aus (9.26) dann keine Basis mehr, und die Dimensionsformel im Allgemeinen falsch.

Definition 11.21. Die Matrix $\mathcal{M}_{CB}(\Phi)$ im Beweis von Theorem 11.20 heisst die *Darstellung* der linearen Abbildung Φ bezüglich den Basen B und C (kurz: *Koordinatendarstellung*, darstellende Matrix oder *Darstellungsmatrix*).

Mnemonik:

$$\text{gegeben: } \Phi(b_i) = c_j \varphi_i^j \leadsto \text{definiere: } \mathcal{M}_{CB}(\Phi) = (\varphi_i^j)_{\substack{j=1, \dots, m \\ i=1, \dots, n}} \quad (11.13)$$

$$\text{betrachte: } \Phi(b_i v^i) = c_j \varphi_i^j v^i = c_j w^j \leadsto \text{folgere: } w^j = \varphi_i^j v^i$$

Kurzform:

$$\Phi(B) = C \cdot \mathcal{M}_{CB}(\Phi), \quad w = \mathcal{M}_{CB}(\Phi) \cdot v \quad (11.14)$$

In Worten: Die Darstellungsmatrix besteht einerseits aus den Koordinaten der Bilder der Basis des Definitionsbereichs bezüglich der Basis des Wertebereichs, und drückt andererseits die Koordinaten im Wertebereich aus durch die Koordinaten im Definitionsbereich.

Aus der Konstruktion (11.12) ist ausserdem ersichtlich: Die Spalten der Darstellungsmatrix sind ein Erzeugendensystem des Bildes von Φ , dargestellt in der Basis des Wertebereichs.

Definition 11.22. Der *Spaltenraum* einer Matrix $\mathcal{A} \in \text{Mat}_{m \times n}(K)$ ist der von den Spalten von $\mathcal{A}$ aufgespannte Unterraum des K^m . Der *Zeilenraum* ist der von den Zeilen von $\mathcal{A}$ aufgespannte Unterraum des K^n (als Vektorraum von Zeilenvektoren).

Der Spaltenraum einer Matrix ist also das Bild der zugehörigen Abbildung $K^n \rightarrow K^m$, Gl. (11.2). Zur Interpretation des Zeilenraums siehe S. 87.

Proposition 11.23. Sind V, W, X drei K -Vektorräume, $\Phi \in \text{Hom}(V, W)$, $\Psi \in \text{Hom}(W, X)$, B eine Basis von V , C eine Basis von W , und D eine Basis von X , so gilt

$$\mathcal{M}_{DB}(\Psi \circ \Phi) = \mathcal{M}_{DC}(\Psi) \cdot \mathcal{M}_{CB}(\Phi) \quad (11.15)$$

Beweis. Ausschreiben: $\Psi \circ \Phi(b_i) = \Psi(c_j \varphi_i^j) = \Psi(c_j) \varphi_i^j = d_k \psi_j^k \varphi_i^j$. $\square$

Lemma 11.24. Sei V ein K -Vektorraum und B, A zwei Basen von V . Dann gilt

$$\mathcal{M}_{BA}(\text{id}_V) = \mathcal{S}_{BA} \quad (11.16)$$

d.h. die Darstellungsmatrix der Identität bezüglich A und B ist gerade die Basiswechselmatrix von B zu A .

Beweis. Folgt unmittelbar aus Definition 10.4 und Definition 11.21, speziell dem Vergleich von (10.11) und (11.13), unter Beachtung von $\text{id}_V(b_i) = b_i$. $\square$

Theorem 11.25 (Basiswechselformel). Ist $\Phi \in \text{Hom}(V, W)$ und sind B, A zwei Basen von V , C, D zwei Basen von W , so gilt

$$\mathcal{M}_{DA}(\Phi) = \mathcal{S}_{DC} \mathcal{M}_{CB}(\Phi) \mathcal{S}_{BA} = (\mathcal{S}_{CD})^{-1} \mathcal{M}_{CB}(\Phi) \mathcal{S}_{BA} \quad (11.17)$$

Beweis. Folgt sofort durch Anwendung von Proposition 11.23 und Lemma 11.24 auf $\text{id}_W \circ \Phi \circ \text{id}_V$. Die letzte Umformung ist eine Konsequenz von Lemma 10.5 (ii): “Zum Wechsel vom Basispaar B/C zum Basispaar A/D benötigt man im Definitionsbe- reich die Basiswechselmatrix und im Wertebereich die Inverse.” $\square$

Beispiel 11.26. Sei $V = K[x]_2$ der 3-dimensionale Vektorraum der Polynome vom Grad ≤ 2 , $W = K^2$, und $\Phi : V \rightarrow W$ gegeben durch die Auswertungsabbildungen an den Stellen 1 und 2,

$$\Phi(p) = \begin{pmatrix} p(1) \\ p(2) \end{pmatrix} \quad (11.18)$$

Bezüglich der Basis $B = (1, x, x^2)$ von V und $C = (e_1, e_2)$ von W ist die Darstellungsmatrix

$$\mathcal{M}_{CB}(\Phi) = \begin{pmatrix} 1 & 1 & 1 \\ 1 & 2 & 4 \end{pmatrix} \quad (11.19)$$

§ 11. LINEARE ABBILDUNGEN

Die Basiswechselmatrix von B zur Basis $A = (1, 2x - 1, 6x^2 - 6x + 1)$ ist⁵⁹

$$\mathcal{S}_{BA} = \begin{pmatrix} 1 & -1 & 1 \\ 0 & 2 & -6 \\ 0 & 0 & 6 \end{pmatrix}, \quad (11.20)$$

die von C zu $D = \left(\begin{pmatrix} 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right)$ ⁶⁰

$$\mathcal{S}_{CD} = \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix} \quad (11.21)$$

Es folgt:

$$\mathcal{M}_{DA}(\Phi) = \begin{pmatrix} 1 & 0 \\ -1 & 1 \end{pmatrix} \begin{pmatrix} 1 & 1 & 1 \\ 1 & 2 & 4 \end{pmatrix} \begin{pmatrix} 1 & -1 & 1 \\ 0 & 2 & -6 \\ 0 & 0 & 6 \end{pmatrix} = \begin{pmatrix} 1 & 1 & 1 \\ 0 & 2 & 12 \end{pmatrix} \quad (11.22)$$

- Es ist ein natürliches Bestreben, die Basen so zu wählen, dass die Darstellungsmatrix besonders einfach wird.

Proposition 11.27. *Es seien V und W endlich-dimensionale K -Vektorräume, und $\Phi \in \text{Hom}(V, W)$. Dann existiert eine Basis B von V und eine Basis C von W so, dass*

$$\mathcal{M}_{CB}(\Phi) = \begin{pmatrix} 1 & \cdots & 0 & 0 & \cdots & 0 \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots \\ 0 & \cdots & 1 & 0 & \cdots & 0 \\ 0 & \cdots & 0 & 0 & \cdots & 0 \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots \\ 0 & \cdots & 0 & 0 & \cdots & 0 \end{pmatrix} = \begin{pmatrix} I_r & 0_{r \times (n-r)} \\ 0_{(m-r) \times r} & 0_{(m-r) \times (n-r)} \end{pmatrix} \in \text{Mat}_{m \times n}(K) \quad (11.23)$$

Hierbei bezeichnet, für alle $k, l \in \mathbb{N}$, $0_{k \times l} \in \text{Mat}_{k \times l}(K)$ die Nullmatrix der angegebenen Grösse. In jeder Darstellung dieser Gestalt gilt $r = \text{rang}(\Phi) = \dim(\text{im}(\Phi))$.

Beweis. 1. Schritt: Sei $d = \text{defect}(\Phi) = \dim(\ker(\Phi))$. Setze $r := n - d$ und wähle eine Basis $(b_{r+1}, \dots, b_n)$ von $\ker(\Phi) \subset V$.

2. Schritt: Ergänze zu einer Basis $(b_1, \dots, b_n)$ von V . Setze $c_i := \Phi(b_i)$ für $i = 1, \dots, r$.

3. Schritt: Behaupte, die Familie $(c_j)_{j=1, \dots, r}$ ist linear unabhängig.

Bew.: Aus

$$\sum_{j=1}^r \beta^j c_j = \Phi \left(\sum_{i=1}^r \beta^i b_i \right) = 0_W \quad (11.24)$$

folgt $\sum_{i=1}^r \beta^i b_i \in \ker(\Phi)$, d.h. es existieren α^i für $i = r+1, \dots, n$ so, dass $\sum_{i=1}^r \beta^i b_i = \sum_{i=r+1}^n \alpha^i b_i$. Da B linear unabhängig ist, folgt (insbesondere) $\beta^j = 0$ für $j = 1, \dots, r$.

⁵⁹Was ist speziell an A ? Die Elemente sind orthogonal bezüglich dem Integral auf dem Intervall $[0, 1]$: Es gilt $\int_0^1 a_i a_j = 0$ falls $i \neq j$.

⁶⁰Hier ist nichts speziell. Zu beachten ist wieder, dass die Koordinaten *kontragredient* zur Basis transformieren: $\Phi(p) = p(1)c_1 + p(2)c_2 = p(1)d_1 + (p(2) - p(1))d_2$.

4. Schritt: Ergänze zu einer Basis $C = (c_1, \dots, c_m)$ von W .

Ergebnis: Die Darstellungsmatrix $\mathcal{M}_{CB}(\Phi)$ hat die angegebene Form. Die ersten r Spalten sind eine Basis von $\text{im}(\Phi)$, dargestellt in der Basis C (siehe Bemerkungen S. 80), d.h. $r = \text{rang}(\Phi)$ (sc. unabhängig davon wie die spezielle Form erzeugt wurde).

Alternativ: Starte mit beliebigen Basen von V und W . Beobachte, dass gemäß Theorem 11.25 bei Basiswechseln vom Typ II, III, IV (siehe S. 58) im Definitionsbereich V die Darstellungsmatrix von rechts mit den Elementarmatrizen (10.20), (10.21), (10.22) multipliziert wird, beziehungsweise bei Basiswechseln im Wertebereich W mit deren Inversen von links. Wir stellen dann zunächst mit dem auf S. 59 beschriebenen Gauss-Algorithmus die holländische Stufenform her. Dies entspricht in etwa dem obigen 2. Schritt. Anschliessend räumen wir mit Hilfe von elementaren Zeilenumformungen alle Spalten unterhalb der Pivot-Elemente frei⁶¹ und normieren zuletzt die Pivot-Elemente auf 1. Dies ergibt wieder die gewünschte Form für die Darstellungsmatrix. $\square$

Korollar 11.28 (Rangsatz). *Es gilt stets*

$$n = \dim(V) = \dim(\ker(\Phi)) + \dim(\text{im}(\Phi)) = \text{defect}(\Phi) + \text{rang}(\Phi) \quad (11.25)$$

Beweis. Abzulesen aus (11.23): Die Basen B und C sind zwar nicht eindeutig, r und d aber sind Dimensionen von Untervektorräumen und damit basisunabhängig. $\square$

Korollar 11.29 (Zeilenrang gleich Spaltenrang). *Die Dimension des Zeilenraums einer Matrix (s. Def. 11.22) ist gleich der Dimension ihres Spaltenraums.*

Beweis. Genauso wie im Beweis von Proposition 11.27 der Spaltenrang ist die Dimension des Zeilenraums invariant unter elementaren Zeilen- und Spaltenumformungen. Sie lässt sich also auch aus Gl. (11.23) ablesen. Beachte: Hier wird nicht benutzt oder ausgesagt (weil es auch nicht richtig ist), dass der Zeilenraum als Unterraum von W oder V interpretiert werden kann. Siehe vielmehr Proposition 12.6. $\square$

Korollar 11.30. *Im Falle $\dim V = \dim W$ sind äquivalent*

- (i) Φ ist ein Isomorphismus.
- (ii) Φ ist injektiv (äquivalent zu $\ker(\Phi) = \{0_V\}$ und $\text{defect}(\Phi) = 0$).
- (iii) Φ ist surjektiv (äquivalent zu $\text{im}(\Phi) = W$ und $\text{rang}(\Phi) = \dim V$).

Beweis. Nicht-trivial ist nur die Äquivalenz von (ii) und (iii). Sie folgt sofort aus Korollar 11.28, und lässt sich mit Hilfe von Korollar 11.16 auch als Konsequenz von Proposition 4.11 begreifen. $\square$

Fazit und Ausblick: Rang (und Defekt) sind die einzigen basisunabhängigen *Invarianten* einer linearen Abbildung zwischen zwei endlich-dimensionalen Vektorräumen. Weitere, interessante Fragestellungen, die teilweise in der HöMa 2 und 3 behandelt werden, ergeben sich unter anderem bei *Einschränkung der erlaubten Basiswechsel*. 1. Häufig interessieren wir uns nicht für lineare Abbildungen und Vektorräume in Isolation, sondern im Wechselspiel mit anderen. Schon *zwei* lineare Abbildungen $\Phi_1, \Phi_2 : V \rightarrow W$ lassen sich im Allgemeinen nicht mehr *gleichzeitig* auf die Gestalt

⁶¹Beachte, dass Spaltenumformungen alleine hierfür im Allgemeinen nicht ausreichen.

§ 12. SUMMEN-, QUOTIENTEN-, DUALRÄUME

(11.23) bringen. Beispiel: $V = \mathbb{R}$, $W = \mathbb{R}$. Erlaubte Basiswechsel: Multiplikation von $e = 1 \in V$ mit $\lambda \in \mathbb{R}^\times$, Multiplikation von $f = 1 \in W$ mit $\mu \in \mathbb{R}^\times$; ändern die Darstellung von $\Phi_i = \varphi_i \in \mathbb{R}$ zu $\mu^{-1}\varphi_i\lambda$, für $i = 1$ und $i = 2$ in gleicher Weise; lassen das *Verhältnis* $\varphi_1 : \varphi_2$ invariant.

2. Im Falle $W = V$ ist es sinnvoll und natürlich, darauf zu bestehen, dass im Definitions- und Wertebereich die gleiche Basis verwendet wird. Dies verhindert im Allgemeinen die Herstellung der Form (11.23), für die wir unabhängige Spalten- und Zeilenumformungen benötigt haben. Beispiel: $V = W = \mathbb{R}^2$, $B = E =$ Standardbasis, $\mathcal{M}_{EE}(\Phi) = \begin{pmatrix} \varphi_1 & 0 \\ 0 & \varphi_2 \end{pmatrix}$. Bei Übergang zur Basis $C = (\lambda_1 e_1, \lambda_2 e_2)$ wird die Multiplikation mit $\begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix}$ von rechts durch die Multiplikation mit ihrer Inversen von links vollständig kompensiert. Die φ_i bleiben unverändert.

3. In der Gegenwart eines euklidischen inneren Produkts (siehe 8.4) möchte man nur Basiswechsel zulassen, die daraus abgeleitete geometrische Größen (Längen, Winkel, Orientierung) unverändert lassen.

§ 12 Summen-, Quotienten-, Dualräume

Mit Hilfe und zur Ergänzung der Erkenntnisse zu linearen Abbildungen wenden wir uns nun der zweiten auf S. 72 aufgeworfenen Frage zu: Wie konstruieren wir “neue” Vektorräume aus “alten”, und welche Beziehungen bestehen zwischen den Ergebnissen aufgrund dieser Konstruktionen?

Ein charakteristisches Beispiel hierfür ist der in Lemma 11.8 eingeführte Vektorraum der linearen Abbildungen. Isoliert betrachtet könnte man $\text{Hom}(V, W)$ einfach mit jedem anderen K -Vektorraum der richtigen Dimension “verwechseln”. Dass dies ein “Irrtum” wäre, erkennt man im Beweis von Lemma 11.20 an der Abhängigkeit des Isomorphismus mit $\text{Mat}_{m \times n}(K)$ von den für die “Bausteine” gewählten Basen B und C . Konkret: Allgemeine Basiswechsel eines $n \cdot m$ -dimensionalen Vektorraums werden beschrieben durch (invertierbare) $nm \times nm$ -dimensionale Matrizen mit $n^2 m^2$ unabhängigen Einträgen. Von V und W stammende Basiswechsel geben höchstens $n^2 + m^2$ Parameter,⁶² das sind im Allgemeinen weniger. Abstrakt (also ohne Basisbezug) entsteht die Diskrepanz genau dadurch, dass $\text{Hom}(V, W)$ eben nicht “für sich alleine steht”, sondern in der “Menge aller Vektorräume” einen speziellen Platz einnimmt und zwar insbesondere stets zusammen mit der “Auswertungsabbildung” $\text{Hom}(V, W) \times V \rightarrow W$, $(\Phi, v) \mapsto \Phi(v)$ geliefert wird. Die Entwicklung des zugehörigen allgemeinen Formalismus würde uns in der HöMa 1 allerdings etwas zu weit führen,⁶³ sodass wir uns hier auf ein paar einfachere (aber durchaus wichtige, und eng verwandte) Konstruktionen beschränken. Die beiden ersten verwenden bereits bekannte Zutaten.

⁶²Tatsächlich ist einer dieser Parameter gar nicht effektiv. Multiplikation von B mit $\alpha \in K^\times$ bei gleichzeitiger Multiplikation von C mit α^{-1} ändert die Darstellungsmatrix nicht.

⁶³Besagter Formalismus ist eng verwandt mit dem für Determinanten, den wir auf S. 65 erwähnt hatten.

Definition 12.1. Für gegebenen K -Vektorraum V heisst der Vektorraum der linearen Abbildungen von V nach K (als ein-dimensionalem K -Vektorraum) der *Dualraum* von V , geschrieben

$$V^\vee = \{\lambda : V \rightarrow K \mid \forall v_1, v_2 \in V, \alpha \in K : \lambda(v_1 + \alpha v_2) = \lambda(v_1) + \alpha \lambda(v_2)\} \quad (12.1)$$

$V^\vee$ ist einerseits linearer Unterraum des Vektorraums $\mathcal{F}(V, K)$ aller K -wertigen Funktionen auf V (als Menge), und andererseits ein Spezialfall von $\text{Hom}(V, W)$ für $W = K$. Der Nullvektor in $V^\vee$ ist die Nullabbildung $0_{V^\vee} : v \mapsto 0 \in K \ \forall v \in V$. Elemente von $V^\vee$ heissen *Linearformen*, *duale Vektoren*, manchmal auch *Kovektoren* oder *lineare Funktionale* (insbesondere im unendlich-dimensionalen Fall).

Proposition 12.2. Sei V ein endlich-dimensionaler Vektorraum. Dann ist $V^\vee$ ebenfalls endlich-dimensional, mit $\dim V^\vee = \dim V$.

Beweis. Ist wegen $\dim K = 1$ (als K -Vektorraum) einfach nur ein Spezialfall von Theorem 11.20, wird aus notationellen Gründen aber hier wiederholt: Wir wählen eine Basis $B = (b_1, \dots, b_n)$ von V und definieren für alle $j = 1, \dots, n$ die Abbildungen $\beta^j \in V^\vee$ als lineare Fortsetzung von

$$\beta^j(b_i) = \delta_i^j = \begin{cases} 1 & \text{falls } j = i \\ 0 & \text{sonst} \end{cases} \quad \text{das heisst also: } \beta^j(v^i b_i) = v^j \quad (12.2)$$

β^j ist dadurch wohldefiniert, da B eine Basis von V ist, vgl. Proposition 11.15 und siehe auch Gl. (11.1).

Beh.: Die Familie $B^\vee = (\beta^1, \dots, \beta^n)$ ist eine Basis von $V^\vee$.

Bew.: $B^\vee$ ist Erzeugendensystem: Für alle $\lambda \in V^\vee$ gilt, $\forall v = v^i b_i \in V$:

$$\lambda(v) = v^i \lambda(b_i) = \lambda(b_j) v^j = \lambda(b_j) \beta^j(v^i b_i) = (\lambda(b_j) \beta^j)(v) \quad (12.3)$$

d.h. $\lambda = \lambda(b_j) \beta^j$.

$B^\vee$ ist linear unabhängig: Angenommen $\lambda_j \beta^j = 0_{V^\vee}$. Dann gilt insbesondere für alle $i = 1, \dots, n$

$$0 = \lambda_j \beta^j(b_i) = \lambda_i \quad (12.4)$$

□

Definition 12.3. Die für gegebene Basis B von V via (12.2) definierte Basis $B^\vee$ von $V^\vee$ heisst die zu B *duale Basis*.

Bemerkungen. 1. Die Hochstellung der Indizes an der dualen Basis (und entsprechende Tiefstellung an den Koordinaten bezüglich $B^\vee$) ist konsistent mit unseren Vereinbarungen für lineare Abbildungen, nach denen eine Basis B von V , zusammen mit der Standardbasis $E = (1)$ von K , den Dualraum $V^\vee = \text{Hom}(V, K)$ mit dem Raum der $1 \times n$ -Matrizen (d.h., Zeilenvektoren der Länge n) identifiziert. Für alle $\lambda = \lambda_j \beta^j \in V^\vee$ ist

$$\mathcal{M}_{EB}(\lambda) = (\lambda_1 \ \cdots \ \lambda_n) \quad (12.5)$$

§ 12. SUMMEN-, QUOTIENTEN-, DUALRÄUME

Diese Notation erlaubt insbesondere die Verwendung von Matrixmultiplikation und Einsteinscher Summenkonvention bei der Auswertung von Linearformen auf Vektoren. Für alle $v = v^i b_i$, $\lambda = \lambda_j \beta^j$ gilt:⁶⁴

$$\lambda(v) = \lambda_j \beta^j(b_i) v^i = \lambda_j \delta_i^j v^i = \lambda_i v^i = (\lambda_1 \cdots \lambda_n) \cdot \begin{pmatrix} v^1 \\ \vdots \\ v^n \end{pmatrix} \quad (12.6)$$

Die Hochstellung der Indizes ist allerdings nicht veträglich mit den Verabredungen über abstrakte “kontext-lose” Vektorräume von Theorem 11.17 und Definition 11.18. Zur Betonung der Unterschiede schreiben wir im Folgenden weiterhin

$$\Phi_{B^\vee} : K^n \rightarrow V^\vee, \quad \begin{pmatrix} \lambda_1 \\ \vdots \\ \lambda_n \end{pmatrix} \mapsto \lambda_j \beta^j \quad (12.7)$$

für die Parametrisierung von $V^\vee$ durch Spaltenvektoren, und (bei Bedarf)

$$\underline{\Phi}_{B^\vee} : \underline{K}^n \rightarrow V^\vee, \quad (\lambda_1, \dots, \lambda_n) \mapsto \sum_{j=1}^n \lambda_j \beta^j = \lambda_j \beta^j \quad (12.8)$$

für die Parametrisierung durch Zeilenvektoren. (Im nächsten Durchlauf der Vorlesung werden wir die Notation $\underline{K}^n$ für $\text{Mat}_{1 \times n}(K)$ bereits in Gl. (9.1) einführen.)

2. Im unendlich-dimensionalen Fall gilt die Aussage von Proposition 12.2 *nicht*.⁶⁶ So hat z.B. der Vektorraum $\mathcal{F}^{\text{fin}}(\mathbb{N}, K)$ der endlichen Folgen die abzählbare Basis $(e_i)_{i \in \mathbb{N}}$ mit

$$e_i = (0, \dots, 0, 1, 0, \dots, 0, \dots) \quad (12.9)$$

$\uparrow$
i-te Stelle

Allerdings sind die zu den e_i dualen Kovektoren $(\epsilon^j)_{j \in \mathbb{N}}$ mit

$$\epsilon^j(e_i) = \delta_i^j \quad (12.10)$$

wie vorher in Gl. (12.2) zwar linear unabhängig, aber kein Erzeugendensystem für den Dualraum $\mathcal{F}^{\text{fin}}(\mathbb{N}, K)^\vee$: Die lineare Abbildung, welche auf der Basis durch $e_i \mapsto 1 \forall i$ erklärt ist, ist keine endliche Linearkombination von $(\epsilon^j)_{j \in \mathbb{N}}$.

3. Aber auch einen endlich-dimensionalen Vektorraum muss man sorgfältig von seinem Dualraum unterscheiden. Für gegebene oder gewählte Basis B von V , mit dualer Basis $B^\vee$ von $V^\vee$, erhält man zwar mit der “tautologischen” Identifikation

⁶⁴Ich empfehle die Vermeidung der (in der mathematischen Literatur weit verbreiteten) Notation $\langle \lambda, v \rangle$ für die “natürliche Paarung” zwischen einem Vektorraum und seinem Dualraum.⁶⁵ Die “Auswertungsnotation” $\lambda(v)$ ist nicht nur ökonomischer, sondern beugt auch der Verwechslung mit allfälligen inneren Produkten vor, die eine viel grössere Gefahr darstellt als die Unlesbarkeit.

⁶⁵Die Notation ist unter anderem nicht symmetrisch zwischen λ und v , es sei denn man setzt sie fort zu $(V \oplus V^\vee) \times (V \oplus V^\vee)$.

⁶⁶Ein Vektorraum ist dann und nur dann isomorph zu seinem Dualraum, wenn er endlich-dimensional ist. Man könnte in Erwägung ziehen, dies als basisfreie Definition zu benutzen.

$K^n = K^n$ einen scheinbar natürlichen Isomorphismus $\Xi_B := \Phi_{B^\vee} \circ \Phi_B^{-1} : V \rightarrow V^\vee$ mit $\Xi_B(b_i) = \beta^i$, d.h., der Darstellungsmatrix

$$\mathcal{M}_{B^\vee B}(\Xi_B) = I_n \quad (12.11)$$

Dieser Isomorphismus *hängt aber von der Wahl der Basis ab* (und ist insofern eben *nicht* natürlich): Ist $C = (c_k)$ mit $c_k = b_i \sigma_k^i$ eine andere Basis von V , und schreiben wir $\gamma^l = \beta^j \tau_j^l$ für die Darstellung der dualen Basis $C^\vee = (\gamma^l)$ als Linearkombinationen von $B^\vee$, so gilt

$$\delta_k^l = \gamma^l(c_k) = \tau_j^l \beta^j(b_i) \sigma_k^i = \tau_j^l \sigma_k^j \quad \begin{array}{l} \text{Zeilenindex} \\ \text{Zeilenindex} \end{array} \quad (12.12)$$

Dies bedeutet für die Basiswechselmatrizen⁶⁷

$$\mathcal{S}_{B^\vee C^\vee} = (\mathcal{S}_{BC})^{-T} = (\mathcal{S}_{BC})^{-T} \quad (12.13)$$

und mit Theorem 11.25 folgt aus (12.11):

$$\mathcal{M}_{C^\vee C}(\Xi_B) = (\mathcal{S}_{B^\vee C^\vee})^{-1} \cdot \mathcal{M}_{B^\vee B}(\Xi_B) \cdot \mathcal{S}_{BC} = (\mathcal{S}_{BC})^T \mathcal{S}_{BC} \quad (12.14)$$

Dies ist im Allgemeinen verschieden von der Darstellungsmatrix des Isomorphismus $\Xi_C = \Phi_{C^\vee} \circ \Phi_C^{-1} : V \rightarrow V^\vee$, der aus Sicht der Basis C “natürlich” erscheint, mit $\mathcal{M}_{C^\vee C}(\Xi_C) = I_n$, d.h. also (im Allgemeinen) $\Xi_C \neq \Xi_B$. *Ohne weitere Struktur gibt es keinen natürlichen Isomorphismus zwischen V und $V^\vee$.* Dies gilt auch für den K^n selbst. Wegen der besonders grossen Verwechslungsgefahr identifizieren wir in diesem Fall gemäss (12.6) konsequent $(K^n)^\vee = \underline{K}^n$.

Zu (12.13) sagt man auch (siehe Fussnote 54): “Kovektoren transformieren unter Basiswechsel kontragredient zu den Vektoren”.

Obwohl anfänglich nicht unbedingt ersichtlich, ist der Dualraum ein essentieller Teil des Formalismus, und spielt auch eine wesentliche Rolle für (angemessen abstrakte) Formulierungen physikalischer Theorien (Mechanik, Elektrodynamik, Relativitätstheorie, Quantenmechanik). Wir benutzen ihn zunächst zur Beantwortung einer offenen Frage von S. 80: Interpretation des Zeilenraums einer Matrix.

Definition 12.4. Sei V ein Vektorraum, und $F \subset V$ eine beliebige Teilmenge. Der *Annihilator* von F ist die Menge aller Linearformen, die auf F verschwinden,

$$\text{ann}(F) := \{ \lambda \in V^\vee \mid \forall f \in F : \lambda(f) = 0 \} \quad (12.15)$$

Lemma 12.5. (i) $\text{ann}(F)$ ist ein Untervektorraum von $V^\vee$.

(ii) $\text{ann}(\emptyset) = V^\vee$, $\text{ann}(V) = \{0_{V^\vee}\}$

(iii) $F_1 \subset F_2 \Rightarrow \text{ann}(F_2) \subset \text{ann}(F_1)$

(iv) $\text{ann}(\text{span}(F)) = \text{ann}(F)$

⁶⁷Beachte: Der Basiswechselformalismus wurde vom Spaltenvektorstandpunkt aus entwickelt, d.h. wir verwenden die Parametrisierung (12.7) mit ihrer “unkonventionellen Indexstellung”: In $\mathcal{S}_{BC} = (\sigma_k^i)$ ist der für die “alte” Basis zuständige Zeilenindex i oben und entsprechend der Spaltenindex k unten. In $\mathcal{S}_{B^\vee C^\vee} = (\tau_j^l)$ ist dies umgekehrt. Der letzte Summationsindex in (12.12) ist in beiden Teilen ein Zeilenindex, d.h. es steht dort wirklich $(\mathcal{S}_{B^\vee C^\vee})^T \mathcal{S}_{BC} = I_n$.

§12. SUMMEN-, QUOTIENTEN-, DUALRÄUME

Beweis. ausschreiben □

Proposition 12.6. Sei $\mathcal{A} = (a_i^j) \in \text{Mat}_{m \times n}(K)$, aufgefasst als lineare Abbildung $K^n \rightarrow K^m$ (und ggfs. als Darstellung einer linearen Abbildung $V \rightarrow W$ bezüglich Basen B, C). Dann gilt:

$$\text{span}(a^1, \dots, a^m) = \text{ann}(\ker(\mathcal{A})) \quad (12.16)$$

In Worten: Der Zeilenraum, aufgefasst als Unterraum von $\underline{K}^n = (K^n)^\vee$ (ggfs. zur Darstellung eines Unterraums von $V^\vee$ mittels $\underline{\Phi}_{B^\vee}$) ist der Annihilator des Kerns, aufgefasst als Unterraum von K^n (ggfs. als Darstellung eines Unterraums von V).

Beweis. (Die Einsichten sind die gleichen wie im Beweis von Proposition 11.27, d.h. eigentlich geht es nur um die Interpretation, und eine explizite Rechnung mit Darstellungsmatrizen.) Wir arbeiten ohne V und W und bezeichnen mit E die Standardbasis von K^n und mit F die von K^m , d.h. also tautologisch $\mathcal{A} = \mathcal{M}_{FE}(\mathcal{A})$. Sei $r = \text{rang}(\mathcal{A})$, $d = n - r = \text{defect}(\mathcal{A})$, und $(b_{r+1}, \dots, b_n)$ eine Basis von $\ker(\mathcal{A})$. Ergänze zu einer Basis $B = (b_1, \dots, b_n)$ von $\mathbb{R}^n$, aufgefasst als Basiswechselmatrix $\mathcal{B} = \mathcal{S}_{EB}$ von der Standardbasis zu B . Dann hat $\mathcal{AB} = \mathcal{M}_{FB}(\mathcal{A})$ die Gestalt

$$\mathcal{AB} = \begin{pmatrix} \tilde{a}_1^1 & \cdots & \tilde{a}_r^1 & 0 & \cdots & 0 \\ \tilde{a}_1^2 & \cdots & \tilde{a}_r^2 & 0 & \cdots & 0 \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots \\ \tilde{a}_1^m & \cdots & \tilde{a}_r^m & 0 & \cdots & 0 \end{pmatrix} \quad (12.17)$$

Die letzten d Spalten sind Null, während der von den Zeilen aufgespannte Unterraum nach wie vor Dimension r hat, siehe Korollar 11.29. Ausserdem gilt für alle $v \in K^n$:

$$v \in \ker \mathcal{A} \Leftrightarrow \mathcal{B}^{-1}v \in \ker \mathcal{AB} \Leftrightarrow \mathcal{B}^{-1}v \text{ ist von der Form } \begin{pmatrix} 0 \\ \vdots \\ 0 \\ \tilde{v}^{r+1} \\ \vdots \\ \tilde{v}^n \end{pmatrix} \quad (12.18)$$

(Das Verschwinden der ersten r Einträge ist äquivalent zur linearen Unabhängigkeit der ersten r Spalten von $\mathcal{AB}$.) Dann gilt für alle $\lambda \in \underline{K}^n$

$$\begin{aligned} \lambda \in \text{ann}(\ker(\mathcal{A})) &\Leftrightarrow \lambda(v) = \lambda \cdot v = 0 \quad \forall v \in \ker \mathcal{A} \subset K^n \\ &\Leftrightarrow \lambda \mathcal{B} \cdot \mathcal{B}^{-1}v = 0 \quad \forall v \in \ker \mathcal{A} \\ &\Leftrightarrow \lambda \mathcal{B} \text{ ist von der Form } (\tilde{\lambda}^1 \cdots \tilde{\lambda}^r 0 \cdots 0) \\ &\Leftrightarrow \lambda \mathcal{B} \text{ ist Linearkombination der Spalten von } \mathcal{AB} \\ &\Leftrightarrow \lambda \text{ ist Linearkombination der Spalten von } \mathcal{A} \end{aligned} \quad (12.19)$$

Dies war die Behauptung. Knobelfrage: Wie sind in diesem Zusammenhang Zeilenumformungen zu interpretieren, und warum werden wir überhaupt mit m Erzeugern von $\text{ann}(\ker(\mathcal{A}))$ konfrontiert? □

- Hier eventuell später: Duale Abbildung, Bidualraum

Definition 12.7. (i) Die *direkte Summe* zweier K -Vektorräume V und W , geschrieben $V \oplus W$, ist das kartesische Produkt $V \times W = \{(v, w) \mid v \in V, w \in W\}$, ausgerüstet mit der komponentenweise Addition und Skalarmultiplikation

$$(v_1, w_1) + (v_2, w_2) = (v_1 + v_2, w_1 + w_2), \quad \alpha \cdot (v, w) = (\alpha v, \alpha w) \quad (12.20)$$

Man schreibt statt $(v, w) \in V \oplus W$ auch $v \oplus w$

(ii) Für $\Phi \in \text{Hom}(V_1, V_2)$ und $\Psi \in \text{Hom}(W_1, W_2)$ definiert man die *direkte Summe von linearen Abbildungen* als

$$\Phi \oplus \Psi : V_1 \oplus W_1 \rightarrow V_2 \oplus W_2, \quad (\Phi \oplus \Psi)(v_1 \oplus w_1) := \Phi(v_1) \oplus \Psi(w_1) \quad (12.21)$$

Lemma 12.8. (i) $0_{V \oplus W} = (0_V, 0_W)$

(ii) $(V \oplus W) \oplus X \cong V \oplus (W \oplus X)$ ⁶⁸

(iii) $V \oplus \{0\} \cong V$

(iv) Ist B eine Basis von V und C eine Basis von W , so ist $B \cup C$ eine Basis von $V \oplus W$. Beachte: Für Familien/geordnete Basen ist eine Anordnung der disjunkten Vereinigung anzugeben.

(v) Sind V und W endlich-dimensional, so gilt $\dim(V \oplus W) = \dim V + \dim W$. (Diese Additivität rechtfertigt im Zweifel die Bezeichnung “direkte Summe”).

(vi) $\text{Hom}(V_1 \oplus W_1, V_2 \oplus W_2) \cong \text{Hom}(V_1, V_2) \oplus \text{Hom}(V_1, W_2) \oplus \text{Hom}(W_1, V_2) \oplus \text{Hom}(W_1, W_2)$.

(vii) Im endlich-dimensionalen Fall induziert die sequentielle Anordnung von Basen eine natürliche Zerlegung der Darstellungsmatrizen

$$\begin{pmatrix} A_{V_1}^{V_2} & A_{W_1}^{V_2} \\ A_{V_1}^{W_2} & A_{W_1}^{W_2} \end{pmatrix}, \quad \text{wo } A_i^j \in \text{Hom}(V_i, W_j) \quad (12.22)$$

(viii) V und W sind Untervektorräume von $V \oplus W$, mit $V \cap W = \{0_{V \oplus W}\}$.

Beweis. alles “offensichtlich”. □

Definition 12.9. Sind U und W Untervektorräume von V , so heisst der von U und W erzeugte Unterraum $U + W = \text{span}(U \cup W) \subset V$ *direkte Summe*, falls $U \cap W = \{0_V\}$, geschrieben $U \oplus W$.

Bemerkungen. Rechtfertigung des Sprachmissbrauchs: In Definition 12.7 (man sagt auch *äussere direkte Summe*) geht es um die Konstruktion eines neuen Vektorraums aus gegebenen. In Definition 12.9 (man sagt auch *innere direkte Summe*) hingegen sind die Verknüpfungen bereits durch die “gemeinsame Umgebung V ” vorgegeben und es geht um die Identifikation zweier Unterräume, die sich trivial schneiden (d.h. nur im Nullvektor). Die abstrakten Eigenschaften sind aber praktisch die gleichen.

⁶⁸Hier und im Weiteren bedeutet $\cong$ “in natürlicher Weise isomorph zu”, ein Begriff, den wir nicht weiter definieren, siehe Fussnote 56.

§ 12. SUMMEN-, QUOTIENTEN-, DUALRÄUME

Definition 12.10. Sei V ein K -Vektorraum.

(i) Untervektorräume U, W von V heissen *komplementär (zueinander)*, falls $U + W = V$ und $U \cap W = \{0\}$, m.a.W., falls $U \oplus W = V$. Man sagt auch: W ist ein *Komplement* zu U (und umgekehrt).

(ii) Unterräume $U_1, \dots, U_N$ heissen *direkte Zerlegung* von V , falls $U_1 + \dots + U_N = V$ und für alle $j = 1, \dots, N$ gilt: $U_j \cap \sum_{i \neq j} U_i = \{0_V\}$, m.a.W. falls $V = \bigoplus_{i=1}^N U_i$.

Beispiel. Ist V endlich-dimensional und $B = (b_1, \dots, b_n)$ eine Basis von V , so gilt

$$V = K \cdot b_1 \oplus K \cdot b_2 \oplus \dots \oplus K \cdot b_n \quad (12.23)$$

wo $K \cdot b_i = \{\alpha b_i \mid \alpha \in K\}$ den von b_i erzeugten ein-dimensionalen Unterraum von V bezeichnet.

Bemerkung. Die Bedingung in Def. 12.10 (ii) stellt sicher, dass V isomorph ist zur “äusseren direkten Summe” der U_i als “abstrakte Vektorräume, bei denen man die Umgebung vergessen hat”. Hierfür reicht es auch aus zu fordern, dass $U_j \cap \sum_{i < j} U_i = \{0_V\}$ für alle $j = 1, \dots, N$. Die Bedingung $U_i \cap U_j = \{0_V\}$ für alle $i \neq j$ ist jedoch zu schwach. Beispiel: Die drei Unterräume $U_1 = \mathbb{R}e_1$, $U_2 = \mathbb{R}e_2$, $U_3 = \mathbb{R}(e_1 + e_2)$ von $\mathbb{R}^2$ haben paarweise trivialen Schnitt, jeder von ihnen ist aber in der direkten Summe der beiden anderen enthalten. $\mathbb{R}^2$ ist nicht direkte Summe von U_1, U_2, U_3 .

Übungsaufgabe. Eine endliche Familie $(U_i)_{i=1, \dots, N}$ von Unterräumen ist genau dann eine direkte Zerlegung von V , wenn jeder Vektor $v \in V$ sich in eindeutiger Weise als Summe $v = \sum_{i=1}^N u_i$ mit $u_i \in U_i$ für alle i schreiben lässt. Ist V endlich-dimensional, so gilt $\dim V = \sum_i \dim U_i$.

Lemma 12.11. Zu jedem UVR U von V existiert ein Komplement W . Ist V endlich-dimensional, so gilt $\dim V = \dim U + \dim W$.

Beweis. Wähle eine Basis A von U , und ergänze sie zu einer Basis B von V . (Dies ist immer möglich nach dem Basisergänzungssatz, Proposition 9.19.)

Beh.: $W = \text{span}(B \setminus A)$ ist ein Komplement zu U .

Bew.: Jedes $v = \sum v^i b_i \in V$ ist eindeutige endliche Linearkombination von B . Die Aufteilung $v = \sum_{i|b_i \in A} v^i b_i + \sum_{i|b_i \notin A} v^i b_i$ zeigt $U + W = V$. Ist andererseits $v \in U \cap W$, so lässt sich v darstellen als $v = \sum_{b_i \in A} v^i b_i$ und als $v = \sum_{b_j \notin A} v^j b_j$. Dann ist $0 = \sum_{b_i \in A} v^i b_i - \sum_{b_j \notin A} v^j b_j$ eine Darstellung des Nullvektors mit paarweise verschiedenen $b_i \neq b_j \forall i \neq j$. Es folgt $v^i = 0 \forall i$.

Die Dimensionsformel ist damit klar, und folgt auch aus Proposition 9.27. $\square$

• Aus der Konstruktion ist ersichtlich, dass je zwei Komplemente eines festen Unterraums isomorph zueinander sind (siehe Korollar 11.16). Als Unterräume von V sind sie jedoch im Allgemeinen verschieden: Basiswechsel von V , die Elementen von $B \setminus A$ Linearkombinationen von A hinzufügen, lassen U fest, verändern aber W .⁶⁹ In der Gegenwart weitere Struktur (innere Produkte, lineare Abbildungen, die U invariant lassen etc.; siehe HöMa 2) lässt sich diese Unbestimmtheit manchmal fixieren. Ansonsten muss man “mit ihr leben”, d.h., sie “zum Prinzip erheben”.

⁶⁹Zugehörige Basiswechselmatrizen haben obere Dreiecksblockform $\mathcal{S} = \begin{pmatrix} I_{|A|} & * \\ 0 & I_{|B \setminus A|} \end{pmatrix}$.

Proposition 12.12. *Sei V ein K -Vektorraum und $U \subset V$ ein Untervektorraum. Dann wird auf V durch*

$$v_1 \sim v_2 \Leftrightarrow v_1 - v_2 \in U \quad (12.24)$$

eine Äquivalenzrelation (siehe Def. 3.5) erklärt. Auf der Menge der Äquivalenzklassen

$$V/\sim = \{[v] := \{\tilde{v} \sim v\} \mid v \in V\} \quad (12.25)$$

(siehe Proposition 4.8) erfüllen die Verknüpfungen

$$[v_1] + [v_2] := [v_1 + v_2], \quad \alpha[v] := [\alpha v] \quad (12.26)$$

die Vektorraumaxiome. Der resultierende Vektorraum heisst Quotientenraum (oder Faktorraum) von V durch (oder nach) U , geschrieben V/U .

Beweis. Äquivalenzrelation: Klar

Wohldefiniertheit: Aus $\tilde{v}_1 - v_1, \tilde{v}_2 - v_2, \tilde{v} - v \in U$ folgt $(\tilde{v}_1 + \tilde{v}_2) - (v_1 + v_2), \alpha\tilde{v} - \alpha v \in U$, da U ein UVR ist. Die Ergebnisse der Verknüpfungen hängen also nicht von den Repräsentanten ab.

Vektorraumaxiome: Ausschreiben bietet keine Schwierigkeiten. Der Nullvektor $0_{V/U} = [0_V] = U$ ist die Äquivalenzklasse aller Vektoren in U .

Notation: Man schreibt für die Elemente des Faktorraums statt $[v]$ äquivalent auch $v + U = \{v + u \mid u \in U\}$ (manchmal sogar nur v). $\square$

Proposition 12.13. *(i) Die kanonische Projektion $\pi : V \rightarrow V/U$, definiert durch $\pi(v) = [v]$, ist ein Epimorphismus mit $\ker(\pi) = U$.*

(ii) Für jedes Komplement W von U ist $\pi|_W : W \rightarrow V/U$ ein Isomorphismus.

(iii) Ist V endlich-dimensional, so gilt $\dim V = \dim U + \dim(V/U)$.

Beweis. (i) Die Surjektivität von π (als mengentheoretische Abbildung) war bereits in Proposition 4.8 gezeigt worden. Die Linearität überprüft man leicht durch Ausschreiben. Die Aussage zum Kern folgt aus $\pi(v) = 0_{V/U} = U \Leftrightarrow v \in U$.

(ii) Aus $V = U + W$ folgt: Für jedes $v \in V$ existieren $u \in U$ und $w \in W$ mit $u + w = v$, so dass $\pi(w) = [w] = [v]$. Also ist π surjektiv. Injektivität folgt wegen $\pi(w) = \pi(\tilde{w}) \Rightarrow w - \tilde{w} \in U \cap W = \{0\}$.

(iii) Folgt aus (ii) und Lemma 12.11. $\square$

Mit anderen Worten ist der Quotientenraum eine Art “Abbild eines idealen universellen Komplements”, dessen Konstruktion allerdings nicht von der Wahl einer Basis abhängt, und in diesem Sinne natürlich ist. Interessant ist vielmehr, dass selbst wenn etwa V “mit einer Basis B gegeben ist”, $\pi(B)$ zwar V/U erzeugt (siehe Lemma 11.12), im Allgemeinen aber nicht mehr linear unabhängig ist. Nur mit diesen Daten hat V/U also definitiv keine ausgezeichnete Basis mehr. (Dies ist auch noch wahr, wenn U mit einer ausgezeichneten Basis A kommt, es sei denn eben, dass $A \subset B$.) Der Zusammenhang lässt sich, ähnlich wie in Proposition 4.8, auch umkehren.

Theorem 12.14 (Homomorphiesatz/Nullter Isomorphiesatz). *Sei $\Phi \in \text{Hom}(V, W)$. Dann ist (in “natürlicher” Weise) $V/\ker(\Phi) \cong \text{im}(\Phi)$.*

§ 13. LINEARE GLEICHUNGEN

Beweis. Jede (mengentheoretische) Abbildung induziert eine surjektive Abbildung auf ihr Bild (triviale Übungsaufgabe). Φ ist ein (linearer) Epimorphismus auf $\text{im}(\Phi)$ (als eigenständiger Vektorraum). Man definiert dann $\phi : V/\ker(\Phi) \rightarrow \text{im}(\Phi)$ durch $\phi([v]) = \Phi(v)$ und überprüft Wohldefiniertheit, Surjektivität und Injektivität wie im Beweis von Proposition 12.13. $\square$

§ 13 Lineare Gleichungen

Den vorläufigen Abschluss bildet in diesem § die “killer app” der linearen Algebra. Problemstellung: Gegeben ein Körper K , natürliche Zahlen $n, m \in \mathbb{N}$, eine *Koeffizientenmatrix* $\mathcal{A} = (a_i^j) \in \text{Mat}_{m \times n}(K)$ und eine *rechte Seite* $b = (b^1, \dots, b^m)^T \in K^m$, beschreibe alle Tupel $x = (x^1, \dots, x^n)^T \in K^n$, für die gleichzeitig alle folgenden Identitäten gelten:

$$\begin{aligned} a_1^1 x^1 + a_2^1 x^2 + \dots + a_n^1 x^n &= b^1 \\ a_1^2 x^1 + a_2^2 x^2 + \dots + a_n^2 x^n &= b^2 \\ &\vdots \\ a_1^m x^1 + a_2^m x^2 + \dots + a_n^m x^n &= b^m, \end{aligned} \tag{13.1}$$

genannt *Lösungen* des *linearen Gleichungssystems* (13.1). Dieses heisst *homogen*, falls $b = 0$, sonst *inhomogen*. b heisst auch *Inhomogenität*. Der *Rang* des Gleichungssystems ist der (Zeilen- oder Spalten-)Rang⁷⁰ der Koeffizientenmatrix.

Idee: Wir interpretieren die in Frage kommenden x als Elemente des n -dimensionalen Vektorraums K^n , die Matrix $\mathcal{A}$ als Darstellung einer linearen Abbildung $\Phi_{\mathcal{A}} : K^n \rightarrow K^m$, nehmen $b \in K^m$, schreiben die Lösungsmenge als

$$\mathcal{L}(\mathcal{A}, b) = \{x \in K^n \mid \mathcal{A} \cdot x = b\}, \tag{13.2}$$

und nutzen unsere Erkenntnisse aus § 11 und § 12.

Beobachtungen:

Proposition 13.1. (i) Für gegebene Koeffizientenmatrix $\mathcal{A}$ ist die Lösungsmenge des homogenen Gleichungssystems $\mathcal{A} \cdot x = 0$ ein linearer Unterraum von K^n der Dimension $d = n - r$, wo r den Rang von (13.1) bezeichnet. Genauer gesagt gilt: $\mathcal{L}(\mathcal{A}, 0) = \ker(\Phi_{\mathcal{A}})$.

(ii) $\mathcal{L}(\mathcal{A}, b)$ ist genau dann nicht leer, wenn $b \in \text{im}(\Phi_{\mathcal{A}})$. In diesem Fall⁷¹ ist $\mathcal{L}(\mathcal{A}, b)$ ein affiner Raum über $\ker(\Phi_{\mathcal{A}})$. Genauer gesagt gilt: Für alle $x_0 \in \mathcal{L}(\mathcal{A}, b)$ ist die Abbildung $\delta(x_0, \cdot) : \mathcal{L}(\mathcal{A}, b) \rightarrow \mathcal{L}(\mathcal{A}, 0)$, $x \mapsto \delta(x_0, x) := x - x_0$ bijektiv und es gilt: $\delta(x_1, x) = \delta(x_1, x_0) + \delta(x_0, x)$ (vgl. Lemma 8.3). In Worten: Existiert überhaupt eine Lösung des inhomogenen Systems, so erhält man jede andere Lösung des inhomogenen Systems durch Addition von genau einer Lösung des zugehörigen homogenen Systems.

⁷⁰Diese sind gleich gemäss Korollar 11.29.

⁷¹Nach unserer Def. 8.2 ist die leere Menge (trivialerweise) ein affiner Raum über jedem beliebigen Vektorraum. Insofern ist die Bedingung nicht wirklich notwendig.

Beweis. (i) Folgt direkt aus der Definition von $\mathcal{L}(\mathcal{A}, 0)$ und $\ker(\Phi_{\mathcal{A}})$. Die Dimensionsformel ist der Rangsatz 11.28.

(ii) Ist $x_0 \in \mathcal{L}(\mathcal{A}, b)$, so ist $b = \Phi_{\mathcal{A}}(x_0) \in \text{im}(\Phi_{\mathcal{A}})$. Die Umkehrung ist eh klar. Für jedes weitere $x \in \mathcal{L}(\mathcal{A}, b)$ ist $\mathcal{A} \cdot (x - x_0) = \mathcal{A} \cdot x - \mathcal{A} \cdot x_0 = 0$, d.h. $x - x_0 \in \mathcal{L}(\mathcal{A}, 0)$; daher ist $\delta(x_0, \cdot)$ wohldefiniert. Injektivität ist klar, die Surjektivität folgt aus $\mathcal{A} \cdot x_0 = b \wedge \mathcal{A} \cdot y = 0 \Rightarrow \mathcal{A} \cdot (x_0 + y) = b$, d.h. $y \in \delta(x_0, \mathcal{L}(\mathcal{A}, b))$.⁷² $\square$

Lösung: Zum Testen von $b \in \text{im}(\Phi_{\mathcal{A}})$ und zur expliziten Parametrisierung von $\ker(\Phi_{\mathcal{A}})$ bringen wir die *erweiterte Koeffizientenmatrix* $(\mathcal{A} | b)$ mit Hilfe von elementaren Zeilenumformungen (in diesem Zusammenhang auch bekannt als Äquivalenzumformungen) auf sog. *reduzierte Zeilenstufenform*,

$$(\mathcal{A} | b) \rightsquigarrow (\tilde{\mathcal{A}} | \tilde{b}) = \left(\begin{array}{cccccccc|cccc|c} 0 & \cdots & 1 & \tilde{a}_{i_1+1}^1 & \cdots & 0 & \cdots & \cdots & 0 & \cdots & \tilde{a}_n^1 & \tilde{b}^1 \\ 0 & \cdots & 0 & 0 & \ddots & 1 & \cdots & \cdots & 0 & \cdots & \tilde{a}_n^2 & \tilde{b}^2 \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \ddots & \ddots & \vdots & \ddots & \vdots & \vdots \\ 0 & \cdots & 0 & 0 & \cdots & 0 & \cdots & \cdots & 1 & \cdots & \tilde{a}_n^r & \tilde{b}^r \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \ddots & \ddots & \vdots & \ddots & \vdots & \vdots \\ 0 & \cdots & 0 & 0 & \cdots & 0 & \cdots & \cdots & 0 & \cdots & 0 & \tilde{b}^m \end{array} \right) \quad (13.3)$$

Im Vergleich zur Transponierten von (10.9) sind die Pivots $\tilde{a}_{i_j}^j = 1$ für $j = 1, \dots, r$ normiert und die einzigen nicht-verschwindenden Einträge in ihrer jeweiligen Spalte.

Ergebnis: Wir können die Lösungsmenge aus (13.3) direkt ablesen.

1. $\mathcal{L}(\mathcal{A}, b) \neq \emptyset \Leftrightarrow \mathcal{L}(\tilde{\mathcal{A}}, \tilde{b}) \neq \emptyset \Leftrightarrow \tilde{b}^{r+1} = \dots = \tilde{b}^m = 0$
2. In diesem Fall ist

$$\begin{aligned} \mathcal{L}(\mathcal{A}, b) &= \mathcal{L}(\tilde{\mathcal{A}}, \tilde{b}) = \{x \mid \\ x^{i_1} &= \tilde{b}^1 - (\tilde{a}_{i_1+1}^1 x^{i_1+1} + \dots + \tilde{a}_{i_2-1}^1 x^{i_2-1}) - \dots - (\tilde{a}_{i_r+1}^1 x^{i_r+1} + \dots + \tilde{a}_n^1 x^n) \\ &\vdots \\ x^{i_r} &= \tilde{b}^r - (\tilde{a}_{i_r+1}^r x^{i_r+1} + \dots + \tilde{a}_n^r x^n) \} \end{aligned} \quad (13.4)$$

Man nennt die $x^{i_1}, \dots, x^{i_r}$, die zu den Pivot-Spalten gehören, auch die “abhängigen”, die übrigen die “unabhängigen” Variablen. (Diese Begriffe sind aber nicht wirklich gut, da die reduzierte Stufenform nicht eindeutig ist.) Knobelaufgaben: (i) Interpretiere den Lösungsalgorithmus in der Basiswechselsprache. (ii) Interpretiere die Inversion von Matrizen “von links” auf S. 63 als simultanes Lösen von linearen Gleichungssystemen.

Beispiel: siehe Übungen

Vermischtes: Für festes $\mathcal{A}$ und variables $b \in \text{im}(\Phi_{\mathcal{A}})$ stehen die Lösungsmengen $\mathcal{L}(\mathcal{A}, b)$ alle in Bijektion zueinander (da sie ja insbesondere in Bijektion zu $\ker(\Phi_{\mathcal{A}}) = \mathcal{L}(\mathcal{A}, 0)$ stehen). Die Bijektionen sind aber nicht kanonisch und als Teilmengen von K^n sind die $\mathcal{L}(\mathcal{A}, b)$ auch verschieden. Per Definition kann jedes $\mathcal{L}(\mathcal{A}, b)$ auch als Element des Quotientenraums $K^n / \ker(\Phi_{\mathcal{A}})$ aufgefasst werden. (Siehe Bemerkung zur

⁷²Die Gültigkeit von $\delta(x_1, x) = \delta(x_1, x_0) + \delta(x_0, x)$ folgt sofort aus der Definition von δ . $\mathcal{L}(\mathcal{A}, b)$ selbst ist aber nicht abgeschlossen unter Vektoraddition in K^n .

§ 13. LINEARE GLEICHUNGEN

Notation im Beweis von Proposition 12.12.) Man nennt $\mathcal{L}(\mathcal{A}, b)$ auch einen *affinen Teilraum* von K^n .

Korollar 13.2. $\mathcal{L}(\mathcal{A}, b) \neq \emptyset \Leftrightarrow \text{rang}(\mathcal{A}) = \text{rang}(\mathcal{A} | b)$ (als $m \times (n+1)$ -dimensionale Matrix).

Beweis. Folgt aus $b \in \text{im}(\Phi_{\mathcal{A}}) \Leftrightarrow b \in \text{span}(a_1, \dots, a_n)$. □

Korollar 13.3. Für eine quadratische Matrix $\mathcal{A} \in \text{Mat}_n(K)$ sind äquivalent:

- (i) Das homogene Gleichungssystem $\mathcal{A} \cdot x = 0$ hat nur die triviale Lösung $x = 0$.
- (ii) Für jedes $b \in K^n$ hat das Gleichungssystem $\mathcal{A} \cdot x = b$ eine eindeutige Lösung $x_* \in K^n$.

(iii) $\mathcal{A}$ ist invertierbar, d.h. $\mathcal{A} \in GL(n, K)$.

In diesem Fall gilt $x_* = \mathcal{A}^{-1}b$, und für seine Komponenten die Cramersche Regel

$$x_*^i = \frac{\det \mathcal{A}_{a_i \leftarrow b}}{\det \mathcal{A}} \quad (13.5)$$

wo man im Zähler in $\mathcal{A} = (a_1, \dots, a_n)$ die i -te Spalte durch b ersetzt hat.

Beweis. Die Äquivalenzen sind gleichbedeutend mit den Äquivalenzen in Korollar 11.30. Gl. (13.5) folgt wegen $b = \mathcal{A} \cdot x_* = a_j x_*^j$ aus

$$\det(a_1, \dots, \underset{\substack{\uparrow \\ i\text{-te Stelle}}}{b}, \dots, a_n) = \det(a_1, \dots, a_j x_*^j, \dots, a_n) = x_*^i \det \mathcal{A} \quad (13.6)$$

wegen Linearität und Alternation der Determinante. Äquivalent: Gemäss Cramerscher Regel 10.22 gilt

$$x_*^i = \frac{1}{\det \mathcal{A}} \sum_{j=1}^n (-1)^{i+j} b^j \det \mathcal{A}_j^i \quad (13.7)$$

Dies ist die Entwicklung nach der i -ten Spalte der Determinante von $\mathcal{A}_{a_i \leftarrow b}$, siehe Proposition 10.19. □

KAPITEL 4

KONVERGENZ

Wir haben Kapitel 3 mit der Einsicht abgeschlossen, dass endlich-dimensionale Vektorräume den natürlichen Rahmen bilden für eine allgemeine Theorie *linearer algebraischer* Gleichungen: Solche Probleme lassen sich stets in endlich vielen Schritten auf die Lösbarkeit von $ax = b$ im zugrundeliegenden Körper zurückführen. Dies ermöglicht die vollständige Beschreibung der Lösungsmenge von Systemen mit beliebig (endlich) vielen Variablen. Das Verständnis der Parametrisierung von Unter- und Faktorräumen erlaubt dann eine zufrieden stellende basisunabhängige geometrisch-physikalische Interpretation.

Andererseits hatten wir in § 6 und § 7 festgehalten, dass gewisse natürlich auftretende nicht-lineare Gleichungen nicht über jedem Körper lösbar sind. Dies kann algebraische Gründe haben (wie bei der Gleichung $x^2 = 2$ über $\mathbb{Q}$) oder durch die Anordnung bedingt sein (wie bei $x^2 = -1$ über $\mathbb{R}$). Ziel dieses Kapitels ist zu illustrieren, wie in der *mathematischen Analysis* eine Vielzahl solcher und verwandter Probleme auf der Grundlage der *Vollständigkeit* im Rechenbereich der reellen oder komplexen Zahlen formuliert und gelöst werden, anstatt durch explizite Algebra oder Vorzeigen. Das unmittelbar nützlichste Ergebnis wird in § 16 zwar “nur” die Definition einiger Zahlen und Funktionen mit Hilfe von Potenzreihen sein. Der intervenierende § 15 legt allerdings die Grundlagen für alle weiteren Untersuchungen, insbesondere für die (mehr-dimensionale) Differential- und Integralrechnung sowie die Theorie der (gewöhnlichen) Differentialgleichungen, die dann in der HöMa 2 und 3 entwickelt werden.

Zunächst aber geht es in § 14 darum, die Charakterisierung der Vollständigkeit von $\mathbb{R}$ über die Supremumseigenschaft,⁷³ die wie bereits angedeutet nicht sehr flexibel ist (und im Komplexen unmittelbar auch keinen Sinn macht), auf verschiedene Weisen umzuformulieren, die sich dann bequem verallgemeinern lassen.

Charakteristisch für die Aussagen der Analysis sind die “für alle $\epsilon > 0$ ”-enthaltenden Wendungen (siehe erstmalig Def. 14.4), und deren in Beweisen häufig benutzte, banale Umformulierungen.

Lemma. Sei $x \in \mathbb{R}$. Dann sind äquivalent:

- (i) $x = 0$
- (ii) Für jedes $\epsilon > 0$ ist $|x| < \epsilon$ (äquivalent dazu: $x \in (-\epsilon, \epsilon)$, dem offenen Intervall).

⁷³Jede nicht-leere nach oben beschränkte Teilmenge von $\mathbb{R}$ besitzt eine kleinste obere Schranke: $\forall T \subset \mathbb{R} : (T \neq \emptyset \wedge (\exists B : \forall x \in T : x \leq B)) \Rightarrow \exists \sup(T) : ((\forall x \in T : x \leq \sup(T)) \wedge (\forall B : \forall x \in T : x \leq B \Rightarrow \sup(T) \leq B))$. Die Eindeutigkeit des Supremums folgt aus der Definition des Begriffs.

(iii) Für jedes $\epsilon > 0$ ist $|x| \leq \epsilon$ (äquivalent dazu: $x \in [-\epsilon, \epsilon]$, dem abgeschlossenen Intervall).

(iv) Es existiert ein $C > 0$ so, dass für jedes $\epsilon > 0$ gilt: $|x| < C \cdot \epsilon$ (bzw. $|x| \leq C \cdot \epsilon$).

(v) Für alle $C > 0$ gilt: Für jedes $\epsilon > 0$ gilt: $|x| < C \cdot \epsilon$ (bzw. $|x| \leq C \cdot \epsilon$).

(vi) Für jede natürliche Zahl $N \in \mathbb{N}$ ist $|x| < \frac{1}{N}$ (bzw. $|x| \leq \frac{1}{N}$, bzw. $\exists C > 0 : \forall N \in \mathbb{N} : |x| < \frac{C}{N}$, bzw. $\forall C > 0 : \forall N \in \mathbb{N} : |x| < \frac{C}{N}$).

Beweis. Dass (i) alle anderen Aussagen impliziert, ist offensichtlich.

(ii) $\Rightarrow$ (i): Wäre $x \neq 0$, so gälte für $\epsilon_* = |x|$: $\epsilon_* > 0$ und $|x| \geq \epsilon_*$. $\nexists$

(iii) $\Rightarrow$ (i): Wäre $x \neq 0$, so gälte für $\epsilon_* = \frac{|x|}{2}$: $\epsilon_* > 0$ und $|x| > \epsilon_*$. $\nexists$

(iv) $\Rightarrow$ (ii): Ist $\epsilon > 0$, so ist auch $\epsilon_{(iv)} = \frac{\epsilon}{C} > 0$. Nach Voraussetzung (iv) folgt $|x| < C \cdot \epsilon_{(iv)} = \epsilon$, die Folgerung von (ii).

(v) $\Rightarrow$ (iv): Offensichtlich mit beliebigem $C > 0$. Die Wahl $C = 1$ impliziert (ii) direkt.

(vi) $\Rightarrow$ (ii): Da $\mathbb{R}$ archimedisch ist (s. Def. 6.7), existiert für jedes $\epsilon > 0$ ein $N_\epsilon \in \mathbb{N}$ so, dass $N_\epsilon > \epsilon^{-1}$, d.h. $\frac{1}{N_\epsilon} < \epsilon$. Aus $|x| < \frac{1}{N}$ für jedes N folgt daraus $|x| < \epsilon$ für besagtes ϵ (und damit jedes $\epsilon > 0$). $\square$

§ 14 Folgen reeller Zahlen

Wir wiederholen und ergänzen einige Vereinbarungen zu $\mathbb{N}$ -indizierten Familien vom Ende von Kapitel 1 und Anfang von Kapitel 3.

Definition 14.1. Sei $X \neq \emptyset$ eine nicht-leere Menge. Unter einer *Folge* in X (oder auch mit Werten in X oder auch X -ige Folge) versteht man eine Abbildung

$$a : \mathbb{N} \rightarrow X \quad n \mapsto a(n) = a_n \in X \quad (14.1)$$

von der Menge der natürlichen Zahlen nach X . Wir schreiben hierfür auch $(a_n)_{n \in \mathbb{N}}$, $(a_n)_{n=1,2,\dots}$, oder (bequem aber unpräzise) $(a_n) \subset X$. Der Wert $a_n \in X$ heisst auch (*n*-tes) *Glied der Folge* (a_n) .

Mit anderen Worten ist eine Folge eine Aufzählung von (wohlgemerkt nicht notwendigerweise verschiedenen) Elementen $a_1, a_2, \dots$ von X . Manche Aussagen über Folgen sind als Aussagen über das Bild, d.h. die Menge $\{a_n \mid n \in \mathbb{N}\} \subset X$ aufzufassen. Andere Aussagen betreffen die eigentliche Abbildung (14.1). Wir interessieren uns dabei insbesondere für den Fall, dass $X = \mathbb{R}$, der bis auf Isomorphismus eindeutige (anordnungs-)vollständige angeordnete Körper.

Definition 14.2. (i) Sind $(x_n)_{n \in \mathbb{N}}$ und $(y_n)_{n \in \mathbb{N}}$ zwei reelle Folgen, so schreiben wir für die (gliedweise) Summe $(x_n + y_n)_{n \in \mathbb{N}}$, das Produkt $(x_n \cdot y_n)_{n \in \mathbb{N}}$, etc.

(ii) Eine Folge (x_n) reeller Zahlen heisst *beschränkt*, wenn eine Schranke $B \geq 0$ existiert so, dass $|x_n| \leq B$ für alle $n \in \mathbb{N}$.

(iii) (x_n) heisst *monoton wachsend*, falls für alle $n, m \in \mathbb{N}$

$$n > m \Rightarrow x_n \geq x_m \quad (14.2)$$

(Dies ist äquivalent zur scheinbar schwächeren Bedingung $x_{n+1} \geq x_n \forall n \in \mathbb{N}$.)

§ 14. FOLGEN REELLER ZAHLEN

(iv) (x_n) heisst *streng monoton wachsend* falls auf der rechten Seite von (14.2) die strenge Ungleichung $x_n > x_m$ gilt.

(v) Entsprechend heisst eine Folge (streng) monoton fallend, falls $n > m \Rightarrow x_n \leq x_m$ (bzw. $x_n < x_m$).

Definition 14.3. Unter einer *Teilfolge* einer Folge $(a_n)_{n \in \mathbb{N}} \subset X$ (für beliebiges X) versteht man eine Folge $(b_k)_{k \in \mathbb{N}}$ in X , welche aus (a_n) durch Angabe einer streng monoton wachsenden Folge $(n_k)_{k \in \mathbb{N}}$ natürlicher Zahlen⁷⁴ via

$$b_k := a_{n_k} \quad (14.3)$$

gewonnen worden ist. Das heisst, man zählt nur einen Teil der Folgenglieder von (a_n) ab, aber unter Berücksichtigung von deren ursprünglichen Reihenfolge.

Beispiele. (i) Die Folge (a_n) mit $a_n = n^2$ ist streng monoton wachsend.

(ii) Die Folge (a_n) mit $a_n = \lfloor \frac{n}{2} \rfloor$, wo für alle $x \in \mathbb{R}$ $\lfloor x \rfloor$ die grösste ganze Zahl kleiner oder gleich x bezeichnet, ist monoton, aber nicht streng monoton, wachsend.

(iii) Die Folge (a_n) mit $a_n = \sin(\pi n/6)$ ist beschränkt. Die Folge (b_k) mit $b_k = a_{6k}$ ist eine Teilfolge von (a_n) mit $b_k = 0$ für alle k .

Die in der Einleitung auf S. 95 angekündigte Strategie zur Lösung analytischer Probleme ist es nun, reelle Zahlen nicht explizit (wie etwa durch Dedekindsche Schnitte) anzugeben, sondern ihnen durch *sukzessive Approximation* mit Folgen allmählich immer näher zu kommen. Dies ist voll an der physikalischen Intuition ausgerichtet.

Definition 14.4. Eine Folge (x_n) reeller Zahlen heisst *konvergent*, falls ein $x \in \mathbb{R}$ existiert mit der Eigenschaft, dass $\forall \epsilon > 0$ ein $N_\epsilon \in \mathbb{N}$ existiert so, dass

$$|x_n - x| < \epsilon \quad \forall n \geq N_\epsilon \quad (14.4)$$

Wir sagen dann auch: (x_n) konvergiert “gegen”, oder “mit Grenzwert” x und schreiben

$$\lim_{n \rightarrow \infty} x_n = x, \text{ manchmal nur } \lim x_n = x, \text{ oder auch } x_n \xrightarrow[n \rightarrow \infty]{} x \quad (14.5)$$

Eine nicht-konvergente Folge heisst *divergent*.

Häufig benutzte, offensichtlich gleichwertige⁷⁵ Varianten der Bedingung (14.4) umfassen:

(i) $\forall \epsilon > 0 : \exists N_\epsilon \in \mathbb{N} : x_n - x \in (-\epsilon, +\epsilon) \quad \forall n \geq N_\epsilon$.

(ii) $\forall \epsilon > 0 : \exists N_\epsilon \in \mathbb{N} : x_n \in (x - \epsilon, x + \epsilon) \quad \forall n \geq N_\epsilon$.

(iii) $\exists C > 0$ so, dass: $\forall \epsilon > 0 : \exists N_\epsilon \in \mathbb{N} : |x_n - x| < C \cdot \epsilon \quad \forall n \geq N_\epsilon$

(iv) Die Benutzung von Worten: Für jedes $\epsilon > 0$ liegen fast alle (d.h. alle bis auf endlich viele) Folgenglieder im Intervall $(x - \epsilon, x + \epsilon)$ oder “weniger als ϵ von x entfernt”.

(v) Die Ersetzung von “ $<$ ” durch “ $\leq$ ”, bzw. der offenen Intervalle durch die abgeschlossenen.

⁷⁴Monotonie für Folgen natürlicher Zahlen ist selbsterklärend genauso definiert wie für reelle Folgen.

⁷⁵Siehe hierzu insbesondere auch das “Triviallemma” auf S. 95.

Beispiele. Eine konstante Folge (x_n) mit $x_n = x$ für alle n konvergiert trivialerweise gegen x . Die Folge (x_n) mit $x_n = x + \frac{1}{n}$ konvergiert gegen x .

Der Umstand, dass zur Entscheidung über die Konvergenz einer Folge der Grenzwert bekannt sein oder erraten werden muss, scheint etwas unzufrieden stellend. Dies wird mit Theorem 14.21 behoben werden. Die *Eindeutigkeit* des Grenzwertes garantiert die Sinnhaftigkeit der Notation (14.5):

Lemma 14.5. *Konvergiert eine reelle Folge (x_n) gegen x , und gegen y , dann gilt notwendiger Weise $x = y$.*

Beweis. Für $\epsilon > 0$ seien N_ϵ und M_ϵ so, dass

$$\begin{aligned} |x_n - x| &< \epsilon \quad \forall n \geq N_\epsilon \\ \text{und } |x_n - y| &< \epsilon \quad \forall n \geq M_\epsilon \end{aligned} \tag{14.6}$$

Dann gilt mit einem beliebigen $n_0 \geq \max\{N_\epsilon, M_\epsilon\}$:

$$|x - y| = |x - x_{n_0} + x_{n_0} - y| \leq |x - x_{n_0}| + |x_{n_0} - y| < 2\epsilon \tag{14.7}$$

Aus diesem Schluss folgt $x = y$. □

Definition 14.6. Eine *Nullfolge* ist eine konvergente Folge (p.t., mit Werten in $\mathbb{R}$) mit Grenzwert 0.

Lemma 14.7. *Für eine reelle Folge (x_n) und $x \in \mathbb{R}$ sind die folgenden Aussagen äquivalent.*

- (i) $\lim_{n \rightarrow \infty} x_n = x$
- (ii) $(x_n - x)_{n \in \mathbb{N}}$ ist eine Nullfolge.
- (iii) $(|x_n - x|)_{n \in \mathbb{N}}$ ist eine Nullfolge.

Beweis. ausschreiben. □

Beispiele und Rechenregeln

Wir entwickeln die wichtigsten Rechenregeln für Konvergenz von Folgen in Konkurrenz mit den folgenden elementaren Grenzwerten.

Beispiel 14.8. (i) Für alle $s \in \mathbb{Q}$ mit $s > 0$ gilt $\lim_{n \rightarrow \infty} \frac{1}{n^s} = 0$.

(ii) Für alle $a > 0$ gilt $\lim_{n \rightarrow \infty} a^{1/n} = 1$.

(iii) $\lim_{n \rightarrow \infty} n^{1/n} = 1$

(iv) Für alle $0 < q < 1$ ist $\lim_{n \rightarrow \infty} q^n = 0$.

(v) Für alle $k \in \mathbb{N}$ und $x > 1$ ist $\lim_{n \rightarrow \infty} \frac{n^k}{x^n} = 0$. “Potenz gewinnt gegen Polynom.”

Der Beweis von (i) beruht auf der Monotonie der Potenzfunktionen, die via Proposition 6.17 (aktuell nur) für alle rationale Exponenten definiert sind: Für jedes $\epsilon > 0$

§ 14. FOLGEN REELLER ZAHLEN

gilt mit $N_\epsilon > \frac{1}{\epsilon^{1/s}}$ (ein solches N_ϵ existiert nach dem Archimedischen Axiom) für alle $n \geq N_\epsilon$:

$$\frac{1}{n^s} \leq \frac{1}{N_\epsilon^s} < (\epsilon^{1/s})^s = \epsilon \quad (14.8)$$

Für (ii) benutzen wir zunächst im Fall $a > 1$ die Bernoullische Ungleichung⁷⁶

$$a = (1 + \underbrace{a^{1/n} - 1}_{>0})^n \geq 1 + n(a^{1/n} - 1) \quad (14.9)$$

um zu folgern, dass für alle n :

$$0 < a^{1/n} - 1 \leq \frac{a - 1}{n} \quad (14.10)$$

Die Aussage folgt dann aus (i) und dem sogenannten “Sandwich-Theorem”:

Lemma 14.9. *Es seien (x_n) eine reelle Folge und (A_n) und (B_n) konvergente reelle Folgen mit*

$$\lim_{n \rightarrow \infty} A_n = \lim_{n \rightarrow \infty} B_n =: x \quad \text{und} \quad A_n \leq x_n \leq B_n \quad \text{für alle } n \quad (14.11)$$

Dann gilt $\lim_{n \rightarrow \infty} x_n = x$.

Beweis. Sind für $\epsilon > 0$ N_ϵ und M_ϵ so, dass $|A_n - x| < \epsilon$ für alle $n \geq N_\epsilon$ und $|B_n - x| < \epsilon$ für alle $n \geq M_\epsilon$, so gilt für alle $n \geq \max\{N_\epsilon, M_\epsilon\}$:

$$-\epsilon < A_n - x \leq x_n - x \leq B_n - x < \epsilon \quad (14.12)$$

und daher $|x_n - x| < \epsilon$. Daraus folgt die Behauptung. $\square$

Bemerkung. · Es genügt offenbar, wenn $x_n \in [A_n, B_n]$ für fast alle n , d.h. $\exists N$ s.d. $x_n \in [A_n, B_n]$ für alle $n \geq N$.

· Wie bereits mehrfach festgehalten, ist für die Abschätzungen im Grossteil der Aussagen und Beweise die Unterscheidung zwischen $<$ und $\leq$ unerheblich.

· Auf der anderen Seite ergeht die *Warnung*, dass beim Grenzübergang selbst Ungleichungen zu Gleichungen werden können. Beispielsweise gilt: Sind (a_n) und (b_n) konvergente Folgen mit $a_n < b_n$ für alle (oder auch “fast alle”) n , so folgt⁷⁷ $\lim a_n \leq \lim b_n$, aber im Allgemeinen gilt nicht “ $<$ ” ! (Konkret: $a_n = 0 \ \forall n$, $b_n = 1/n$.)

Proposition 14.10. *Für reelle Folgen (x_n) und (y_n) gelte $x_n \rightarrow x$ und $y_n \rightarrow y$. Dann gilt*

$$(\alpha) \lim(x_n + y_n) = x + y$$

$$(\beta) \lim(x_n \cdot y_n) = x \cdot y$$

$$(\gamma) \lim |x_n| = |x|$$

$$(\delta) \text{ falls } y \neq 0, \text{ so sind fast alle } y_n \neq 0 \text{ und } \lim \frac{x_n}{y_n} = \frac{x}{y}$$

⁷⁶Siehe Lemma 6.3: Für alle $x \geq -1$, $n \in \mathbb{N}$, gilt $(1+x)^n \geq 1+nx$.

⁷⁷ Beweis: $\lim a_n < a_n + \epsilon < b_n + \epsilon < \lim b_n + 2\epsilon$ für n gross genug, für alle $\epsilon > 0$.

Beweis. siehe Lehrbücher. $\square$

· Natürlich gelten hier keine Umkehrungen. Insbesondere folgt aus $|x_n| \rightarrow |x|$ *nicht* $x_n \rightarrow x$, siehe etwa Beispiel 14.12.

Die Aussage des Beispiels (ii) im Falle $a < 1$ folgt nun aus dem bereits bewiesenen Fall $a > 1$ zusammen mit der Regel (δ) . (Für $a = 1$ ist die Aussage sowieso trivial.) Zum Beweis von (iii) wenden wir die Bernoullische Ungleichung zunächst an auf

$$n^{1/2} = (1 + n^{1/2n} - 1)^n \geq 1 + n(n^{1/2n} - 1) \quad (14.13)$$

$$\text{um zu folgern, dass } 0 \leq n^{1/2n} - 1 \leq \frac{n^{1/2} - 1}{n} = \frac{1}{n^{1/2}} - \frac{1}{n} \xrightarrow{n \rightarrow \infty} 0$$

Mit Lemma 14.9 und Lemma 14.7 folgt daraus $\lim_{n \rightarrow \infty} n^{1/2n} = 1$, und wegen Regel (β)

$$\lim_{n \rightarrow \infty} n^{1/n} = \lim_{n \rightarrow \infty} (n^{1/2n} \cdot n^{1/2n}) = 1 \quad (14.14)$$

· Beispiel (iv) ist genau die Aussage von Prop. 6.11 (die wir auch mit der Bernoullischen Ungleichung bewiesen hatten). Zum Beweis von (v) folgern wir zunächst (noch einmal) aus der Bernoullischen Ungleichung, dass für alle $y > 1$

$$y^n = (1 + y - 1)^n \geq 1 + n(y - 1) > n^{1/2} \cdot n^{1/2} \underbrace{(y - 1)}_{>0} \quad (14.15)$$

$$\text{das heisst also: } \frac{n^{1/2}}{y^n} < \frac{1}{n^{1/2}(y - 1)} \xrightarrow{n \rightarrow \infty} 0$$

Nun schliessen wir mit $y = x^{1/2k} > 1$ und wiederholter Anwendung der Regel (β) :

$$\frac{n^k}{x^n} = \underbrace{\left(\frac{n^{1/2}}{(x^{1/2k})^n} \right)^{2k}}_{\xrightarrow{n \rightarrow \infty} 0} \xrightarrow{n \rightarrow \infty} 0 \quad (14.16)$$

· Wir bemerken, dass im Vergleich zur linearen Algebra die Beweise der Analysis häufig von geschickten Umformungen und von frühestens im nachhinein offensichtlichen Abschätzungen abhängen.

· Ausserdem sind die Folgen (i)–(v), sowie die Rechenregeln, zwar sehr nützlich, illustrieren aber gerade *nicht* die wesentliche Idee, Folgen zum Herzeigen “neuer” reeller Zahlen (insbesondere solcher in $\mathbb{R} \setminus \mathbb{Q}$) zu benutzen. Dafür müssen wir aus der Vollständigkeit von $\mathbb{R}$ Kriterien entwickeln, die die Konvergenz von gewissen Folgen garantieren, gerade auch ohne, dass wir den Grenzwert a priori explizit kennen. (Vielmehr wird dann die reelle Zahl als der Grenzwert der fraglichen Folge definiert.)

· Unser Paradebeispiel hierfür ist die rekursiv definierte Folge

$$x_{n+1} = x_n - \frac{x_n^2 - 2}{2x_n} = \frac{x_n}{2} + \frac{1}{x_n} \quad (14.17)$$

(vgl. (6.15)), welche für beliebigen Startwert $x_0 > \sqrt{2}$ die Eigenschaft hat, dass

$$\lim_{n \rightarrow \infty} x_n = \sqrt{2}, \quad (14.18)$$

was wir als Konsequenz aus unserem ersten grenzwertfreien Konvergenzkriterium bald beweisen. Zunächst noch eine kleine Übung zu Teilfolgen.

§ 14. FOLGEN REELLER ZAHLEN

Lemma 14.11. Eine Folge $(x_n)_{n \in \mathbb{N}}$ konvergiert genau dann, wenn jede ihrer Teilfolgen $(x_{n_k})_{k \in \mathbb{N}}$ konvergiert. Es gilt dann

$$\lim_{k \rightarrow \infty} x_{n_k} = \lim_{n \rightarrow \infty} x_n \quad (14.19)$$

für jede solche Teilfolge.

Beweis. Angenommen, (x_n) konvergiert. Sei $x := \lim_{n \rightarrow \infty} x_n$ ihr Grenzwert, und (x_{n_k}) eine Teilfolge. Für beliebiges $\epsilon > 0$ sei dann N_ϵ so gross, dass $|x_n - x| < \epsilon$ für alle $n \geq N_\epsilon$. Wegen der Monotonie der die Teilfolge parametrisierenden Folge (n_k) gilt $n_k \geq k \ \forall k$, und damit folgt $|x_{n_k} - x| < \epsilon$ für alle $k \geq N_\epsilon$. Die Teilfolge konvergiert also auch gegen x .

· Nehmen wir umgekehrt an, dass jede Teilfolge von (x_n) konvergiert, so folgt die Konvergenz von (x_n) bereits aus der Tatsache, dass (x_n) trivialerweise eine Teilfolge ihrer selbst ist. Damit folgt auch die Aussage zum Grenzwert aller (anderen) Teilfolgen. $\square$

Beispiele 14.12. Es genügt natürlich nicht, wenn irgendeine Teilfolge konvergiert: Die Folge $(a_n)_{n \in \mathbb{N}} = (-1, 1, -1, \dots)$ mit

$$a_n := (-1)^n = \begin{cases} -1 & n \text{ ungerade} \\ +1 & n \text{ gerade} \end{cases} \quad (14.20)$$

ist divergent. Die Teilfolge der Glieder mit geradem Index, $(b_k)_{k \in \mathbb{N}} = (a_{2k})_{k \in \mathbb{N}}$ (d.h. $n_k = 2k$, $b_k = a_{2k}$) ist hingegen konvergent (da konstant gleich $+1$) mit Grenzwert 1 . Ebenso ist $(a_{2k-1})_{k \in \mathbb{N}}$ konvergent mit Grenzwert -1 .

· Es genügt nicht einmal, wenn “in einer Familie von Teilfolgen, welche zusammen alle Folgenglieder abdecken, alle gegen den gleichen Grenzwert konvergieren”. Sei zum Beispiel $a_n = \frac{1}{k}$ wenn n von der Form $n = p^k$ ist für eine Primzahl p , und $a_n = 0$ sonst (wenn also n keine Primzahlpotenz ist). Dann ist für jede Primzahl p die Teilfolge $(a_{n_k})_{k \in \mathbb{N}}$ mit $n_k = p^k$ eine Nullfolge, ebenso wie die komplementäre Teilfolge, die über die n mit mehr als einem Primzahlfaktor gebildet wird (und die identisch 0 ist). Andererseits existiert für jede (noch so grosse) natürliche Zahl N eine Primzahl $p > N$, für die dann $a_p = 1$. Die Teilfolge $(a_p)_{p \text{ prim}}$ ist also konstant gleich 1 und (a_n) selbst ist divergent.

· Es ist instruktiv sich zu überlegen, dass wir für die Untersuchung des Konvergenzverhaltens einer Folge stets endlich viele Folgenglieder ignorieren können. Als Beispiel hierfür betrachten wir für $x \in \mathbb{R}$ die Folge

$$a_n = \frac{x^n}{n!} \quad (14.21)$$

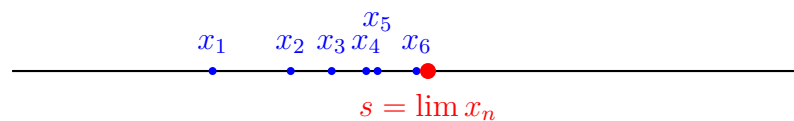
Wir behaupten, (a_n) ist für jedes solche x eine Nullfolge. Sei dazu (für festes x) N_0 so gross, dass $\frac{|x|}{N_0} < \frac{1}{2}$. Dann ist für alle $n \geq N_0$:

$$\begin{aligned} |a_n| &= \left| a_{N_0} \cdot \frac{x^{n-N_0}}{(N_0+1) \cdots (n-1) \cdot n} \right| \\ &= |a_{N_0}| \cdot \frac{|x|}{N_0+1} \cdots \frac{|x|}{n-1} \cdot \frac{|x|}{n} \\ &\leq |a_{N_0}| \cdot 2^{N_0-n} \rightarrow 0 \quad \text{für } n \rightarrow \infty \end{aligned} \quad (14.22)$$

Folgen und Vollständigkeit

Konvergenz lässt sich über den Absolutbetrag prinzipiell in jedem angeordneten Körper definieren, und die obigen Rechenregeln sowie die meisten Beispiele gelten dann fort. Auf den folgenden Seiten wird gezeigt, dass die *Supremumseigenschaft* von $\mathbb{R}$ dreierlei Konvergenzkriterien von Folgen nach sich zieht: Das *Prinzip der monotonen Konvergenz*, das *Intervallschachtelungsprinzip* (welches den *Satz von Bolzano-Weierstrass* impliziert), und das *Cauchy-Kriterium*. Tatsächlich ist jedes davon (gegebenenfalls unter der Voraussetzung des archimedischen Axioms) äquivalent zur Anordnungsvollständigkeit und erlaubt damit alternative Charakterisierungen der reellen Zahlen.

Theorem 14.13 (Prinzip der monotonen Konvergenz). *Jede monoton wachsende und nach oben beschränkte Folge reeller Zahlen ist konvergent.*



Beweis. Sei (x_n) eine monoton wachsende nach oben beschränkte reelle Folge. Die Bildmenge $T = \{x_n \mid n \in \mathbb{N}\} \subset \mathbb{R}$ ist dann nichtleer und nach oben beschränkt. Nach der Supremumseigenschaft (Anordnungsvollständigkeit, Definition 6.14) existiert eine kleinste obere Schranke,

$$s := \sup(\{x_n \mid n \in \mathbb{N}\}) \quad (14.23)$$

Beh.: $\lim_{n \rightarrow \infty} x_n = s$

Bew.: Für jedes $\epsilon > 0$ ist $s - \epsilon < s$, daher keine obere Schranke für T . Es existiert also ein N_ϵ mit $x_{N_\epsilon} > s - \epsilon$. Wegen der Monotonie gilt dann für all $n \geq N_\epsilon$

$$\begin{array}{c} (x_n) \text{ ist monoton wachsend} \\ \downarrow \\ s \geq x_n \geq x_{N_\epsilon} > s - \epsilon \\ \uparrow \\ s \text{ ist obere Schranke} \end{array} \quad (14.24)$$

d.h. $|x_n - s| < \epsilon$, $\forall n \geq N_\epsilon$. Dies impliziert die Behauptung. $\square$

Genauso ist jede monoton fallende und nach unten beschränkte reelle Folge konvergent. Es ist (insbesondere im Kontext der Integralrechnung) zweckmässig, die Begriffe von Supremum/Infimum und Konvergenz so zu erweitern, dass die Supremumseigenschaft und das Prinzip der monotonen Konvergenz auch ohne die Bedingungen “nicht-leer” und “beschränkt” gelten.

Definition 14.14. (i) Ist $T \subset \mathbb{R}$ nicht nach oben (unten) beschränkt, so setzt man $\sup T = +\infty$ (bzw. $\inf T = -\infty$).

(ii) Wir setzen: $\sup \emptyset = -\infty$, $\inf \emptyset = +\infty$.⁷⁸

⁷⁸Dies stellt zusammen mit (i) sicher, dass für alle $T \subset K$ stets $\sup T = \inf\{B \mid x \leq B \ \forall x \in T\}$ (bzw. $\inf T = \sup\{B \mid x \geq B \ \forall x \in T\}$), ist aber nicht mehr verträglich mit $\inf T \leq \sup T$.

§ 14. FOLGEN REELLER ZAHLEN

(iii) Eine Folge reeller Zahlen (x_n) heisst *uneigentlich konvergent gegen $+\infty$ (bzw. $-\infty$)* (man sagt auch: *bestimmt divergent*), falls für alle $B \in \mathbb{R}$ ein N_B existiert so, dass $x_n > B$ für alle $n \geq N_B$.

(iv) Die Menge $\overline{\mathbb{R}} := \mathbb{R} \cup \{-\infty, \infty\}$ heisst *erweiterte Zahlengerade*.

Beispiele. Für jedes $x > 1$ divergiert die Folge

$$a_n = \frac{x^{n^2}}{n!} \quad (14.25)$$

bestimmt gegen $+\infty$. (Übungsaufgabe)

• Jede monotone Folge reeller Zahlen ist entweder konvergent oder bestimmt divergent.

Prinzip der monotonen Konvergenz $\Rightarrow$ Supremumseigenschaft. (Wird nicht in der Vorlesung vorgeführt; alternativ leitet man aus der monotonen Konvergenz zunächst das Intervallschachtelungsprinzip her und daraus die Supremumseigenschaft.)

Proposition. *Sei K ein angeordneter Körper mit der Eigenschaft, dass jede monoton wachsende nach oben beschränkte Folge konvergiert. Dann hat jede nicht-leere nach oben beschränkte Menge eine kleinste obere Schranke.*

Beweis. Wir zeigen zunächst, dass K archimedisch ist: Wäre $\mathbb{N} \subset K$ nach oben beschränkt, so wäre $(n)_{n \in \mathbb{N}}$ eine monoton wachsende nach oben beschränkte Folge. Ist $x = \lim_{n \rightarrow \infty} n \in K$ ihr Grenzwert so konvergiert die Folge $(n+1)_{n \in \mathbb{N}}$ als Teilfolge von $(n)_{n \in \mathbb{N}}$ aufgrund von Lemma 14.11 ebenfalls gegen x , aufgrund der Rechenregel (α) aber gegen $x+1$. Dies ist wegen Lemma 14.5 unmöglich. $\nexists$

Sei nun $T \subset K$ nicht-leer und nach oben beschränkt, und $U = \{y \in K \mid x \leq y \ \forall x \in T\} \subset K$ die Menge der oberen Schranken von T . U ist nicht-leer und nach oben ordnungsabgeschlossen, d.h. für alle $y \in U$, $y' \in K$ mit $y' > y$ gilt auch $y' \in U$. Insbesondere gilt: $x \notin U \Rightarrow x < y \ \forall y \in U$.

Beh.: Für alle $\epsilon > 0$ gilt

$$U \subset U - \epsilon = \{x \mid x + \epsilon \in U\} \not\subset U \quad (14.26)$$

Bew.: Die erste Inklusion folgt aus der Ordnungsabgeschlossenheit, die zweite durch Widerspruch: Aus $U - \epsilon \subset U$ folgt rekursiv $U - n\epsilon \subset U - (n-1)\epsilon \subset \dots \subset U$ für alle $n \in \mathbb{N}$, und daraus aus der archimedischen Eigenschaft (und der Tatsache, dass $U \neq \emptyset$), dass $x \in U$ für alle $x \in K$. Andererseits gilt für beliebiges $x \in T$ (T ist auch nicht-leer) $x - 1 \notin U$. $\nexists$

Beh.: Für alle $n \in \mathbb{N}$ gilt $U - \frac{1}{n} \supset U - \frac{1}{n+1} \supset \dots \supset U$ und

$$\bigcap_{n \in \mathbb{N}} \left(U - \frac{1}{n} \right) \subset U \quad (14.27)$$

Bew.: Die Inklusionskette folgt wieder aus der Ordnungsabgeschlossenheit. Für (14.27) zeigen wir, dass jedes $y \in \bigcap_{n \in \mathbb{N}} (U - \frac{1}{n})$ obere Schranke von T ist: Andernfalls existiert ein $x \in T$ mit $x > y$. Da K archimedisch ist, existiert dann ein $n_0 \in \mathbb{N}$ so, dass auch noch $y + \frac{1}{n_0} < x$. Wegen $y + \frac{1}{n_0} \in U$ muss aber $x \leq y + \frac{1}{n_0}$ gelten. $\nexists$

Setze nun $n_1 = 1$ und wähle $x_1 \in (U - \frac{1}{n_1}) \setminus U$ (diese Menge ist nicht-leer wegen (14.26)). Da $x_1 \notin U$ existiert wegen (14.27) ein $n_2 > n_1$ so, dass $x_1 \notin U - \frac{1}{n_2}$. Wähle $x_2 \in (U - \frac{1}{n_2}) \setminus U$ (ebenfalls nicht-leer wegen (14.26)) und $n_3 > n_2$ so, dass $x_2 \notin U - \frac{1}{n_3}$, etc. Dies liefert eine monoton wachsende Folge ganzer Zahlen $(n_k)_{k \in \mathbb{N}}$ und eine Folge $(x_k)_{k \in \mathbb{N}}$ in K so, dass $x_k \in (U - \frac{1}{n_k}) \setminus (U - \frac{1}{n_{k+1}}) \subset (U - \frac{1}{n_k}) \setminus U$ für alle k . (x_k) ist dann monoton wachsend und durch jedes $y \in U$ nach oben beschränkt. Sie konvergiert nach Voraussetzung und ihr Grenzwert $s = \lim_{k \rightarrow \infty} x_k$ ist nach Fussnote 77 auch kleiner oder gleich als jedes Element von U . Wegen $s \geq x_k$ für alle k in Verbindung mit (14.27) gilt $s \in \bigcap_{k \in \mathbb{N}} (U - \frac{1}{n_k}) = \bigcap_{n \in \mathbb{N}} (U - \frac{1}{n}) \subset U$. s ist also obere Schranke von T und keine obere Schranke ist kleiner. Dies zeigt die Supremumseigenschaft. $\square$

Mit Theorem 14.13 sind wir nun in der Lage, die auf S. 29 geöffnete Klammer endlich zu schliessen.

Beweis von (14.18). Aus den Betrachtungen um Gl. (6.16) folgt, dass die für beliebigen Startwert $x_0 > \sqrt{2}$ via $x_{n+1} = \frac{x_n}{2} + \frac{1}{x_n}$ definierte Folge (x_n) monoton fällt und nach unten (durch $\sqrt{2}$) beschränkt ist. Nach Theorem 14.13 existiert $x = \lim_{n \rightarrow \infty} x_n$, und nach den obigen Rechenregeln gilt

- $x > 0$ (denn $x_n > \sqrt{2} \forall n$)
- Die Folge (y_n) mit $y_n := x_{n+1}$ konvergiert als Teilfolge von (x_n) ebenfalls, mit dem gleichen Grenzwert x (s. 14.11).

Andererseits gilt wegen der Rechenregeln (α) , (δ) :

$$x = \lim_{n \rightarrow \infty} y_n = \lim_{n \rightarrow \infty} \frac{x_n}{2} + \lim_{n \rightarrow \infty} \frac{1}{x_n} = \frac{x}{2} + \frac{1}{x} \quad (14.28)$$

und deshalb $x = \sqrt{2}$

$\square$

Dieses insbesondere von Newton verfochtene Verfahren lässt sich offensichtlich zur Berechnung aller in Proposition 6.17 angekündigten Wurzeln verallgemeinern. Zum Beweis des Fundamentalsatzes der Algebra dringen wir damit aber nicht durch, da $\mathbb{C}$ nicht angeordnet ist. Eine alternative Idee ist die “sukzessive Eingrenzung” von (gesuchten, reellen) Zahlen auf immer kleiner werdende beschränkte Teilmengen.

Definition 14.15. Eine *Intervallschachtelung* ist eine Folge von abgeschlossenen und beschränkten Intervallen $([a_n, b_n])_{n \in \mathbb{N}}$ mit

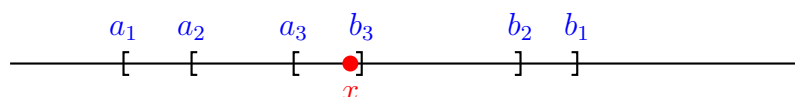
$$\begin{aligned} [a_{n+1}, b_{n+1}] &\subset [a_n, b_n] && \forall n \\ \text{und } \lim_{n \rightarrow \infty} (b_n - a_n) &= \lim_{n \rightarrow \infty} |b_n - a_n| = 0 \end{aligned} \quad (14.29)$$

(Zur Erinnerung: Abgeschlossene und beschränkte Intervalle heissen auch *kompakt*. Ein-Punkt-Mengen sind nach unserer Definition 6.16 erlaubt.)

§ 14. FOLGEN REELLER ZAHLEN

Theorem 14.16 (Intervallschachtelungsprinzip). *Zu jeder Intervallschachtelung $([a_n, b_n])$ existiert genau eine reelle Zahl x , die in allen $[a_n, b_n]$ enthalten ist, d.h.*

$$\bigcap_{n=1}^{\infty} [a_n, b_n] = \{x\} \quad (14.30)$$



Beweis. Die Folge (a_n) ist monoton wachsend und nach oben beschränkt (durch irgendein b_n), daher konvergent wegen Theorem 14.13. Wir setzen $x = \lim a_n$, so dass gilt $x \geq a_n \forall n$. Für gegebenes $\epsilon > 0$ sei nun N_ϵ derart, dass gleichzeitig

$$\text{und } \left. \begin{array}{l} b_n - a_n < \epsilon \\ x - a_n < \epsilon \end{array} \right\} \forall n \geq N_\epsilon \quad (14.31)$$

(Ersteres lässt sich erreichen wegen $\lim(b_n - a_n) = 0$, letzteres wegen $\lim a_n = x$.) Dann gilt $\forall n \geq N_\epsilon$

$$\begin{array}{c} a_n \leq b_n \\ \downarrow \\ -\epsilon < a_n - x \leq b_n - x \leq b_n - a_n < \epsilon \end{array} \quad \begin{array}{c} (14.31) \\ \swarrow \\ \downarrow \\ x \geq a_n \end{array} \quad (14.32)$$

Es folgt $\lim b_n = x$, und wegen der Monotonie von (b_n) , dass $x \leq b_n$, also $x \in [a_n, b_n]$, für alle n . Für jedes (weitere) $y \in \bigcap_n [a_n, b_n]$ folgt aus $\lim |b_n - a_n| = 0$, dass $|x - y| < \epsilon$ für alle $\epsilon > 0$ und daher $y = x$. $\square$

Das Prinzip gilt *nicht* mit offenen Intervallen, z.B. ist $\bigcap_{n=1}^{\infty} (0, \frac{1}{n}) = \emptyset$.

Intervallschachtelungen lassen sich bereits sinnvoll auf nicht angeordnete Rechenbereiche verallgemeinern (siehe Proposition 15.24 in § 15). Wir benutzen sie hier wie angekündigt zum tieferen Verständnis des Konvergenzverhaltens von Folgen.

Theorem 14.17 (Satz von Bolzano-Weierstrass). *Jede (nach oben und unten) beschränkte reelle Folge (x_n) besitzt eine konvergente Teilfolge.*

Beweis. Da (x_n) beschränkt ist, existiert ein $B > 0$ so, dass

$$\{x_n \mid n \in \mathbb{N}\} \subset [-B, B] \quad (14.33)$$

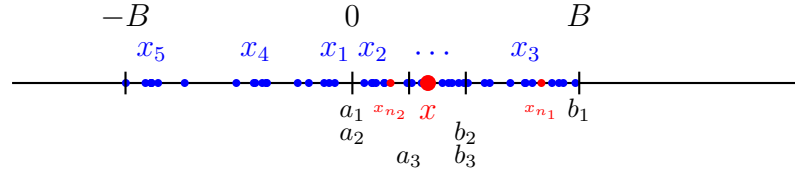
Zur Identifikation eines mutmasslichen Grenzwertes definieren wir rekursiv eine Intervallschachtelung $([a_k, b_k])_{k \in \mathbb{N}}$ so, dass für jedes k gilt:

$$\text{Bed}(k) : [a_k, b_k] \text{ enthält unendlich viele } x_n \quad (14.34)$$

Dazu der 1. Schritt: Falls in $[0, B]$ unendlich viele x_n liegen, setze $[a_1, b_1] := [0, B]$. Andernfalls liegen in $[-B, 0]$ unendlich viele x_n und wir setzen $[a_1, b_1] := [-B, 0]$.

k -ter Schritt: Sei $M := \frac{a_{k-1}+b_{k-1}}{2}$. Da $[a_{k-1}, b_{k-1}]$ wegen $\text{Bed}(k-1)$ unendlich viele x_n enthält, enthält mindestens eines von $[a_{k-1}, M]$ oder $[M, b_{k-1}]$ unendlich viele. Wir setzen $[a_k, b_k] = [M, b_{k-1}]$ falls dieses Intervall unendlich viele x_n enthält, sonst $[a_k, b_k] = [a_{k-1}, M]$. Damit ist $\text{Bed}(k)$ erfüllt und wir können fortfahren.

Per Konstruktion gilt $[a_{k+1}, b_{k+1}] \subset [a_k, b_k]$ und $b_k - a_k = \frac{B}{2^{k-1}} \xrightarrow{k \rightarrow \infty} 0$. Es liegt also eine Intervallschachtelung vor. Sei $x \in \mathbb{R}$ ihr gemäss Theorem 14.16 eindeutige Durchschnitt.



Zur Konstruktion einer Teilfolge $(x_{n_k})_{k \in \mathbb{N}}$ mit Grenzwert x sei n_1 ein Index so, dass $x_{n_1} \in [a_1, b_1]$ und dann rekursiv n_k ein Index, grösser als n_{k-1} so, dass $x_{n_k} \in [a_k, b_k]$. (Der Witz ist, dass, falls es kein solches n_k gäbe, dann lägen höchstens endlich viele x_n in $[a_k, b_k]$, nämlich höchstens alle mit $n \leq n_{k-1}$, im Widerspruch zu $\text{Bed}(k)$.)

Beh.: $\lim_{k \rightarrow \infty} x_{n_k} = x$

Bew.: Für $\epsilon > 0$ sei $K \in \mathbb{N}$ so, dass $\frac{B}{2^{K-1}} < \epsilon \forall k \geq K$. Es gilt $x_{n_k} \in [a_k, b_k] \subset [a_K, b_K] \forall k \geq K$, ausserdem ist per Definition $x \in [a_K, b_K]$. Wegen $b_K - a_K < \epsilon$ folgt $|x - x_{n_k}| < \epsilon \forall k \geq K$. $\square$

In diesem Beweis haben wir bei der Konstruktion der Intervallschachtelung stets das obere Halbintervall vorgezogen. Würden wir im Falle, dass $[a_{k-1}, M]$ und $[M, b_{k-1}]$ beide unendlich viele Folgenglieder enthalten, auch mal das untere auswählen, so hätte die resultierende Teilfolge (im Allgemeinen) einen anderen Grenzwert.

Definition 14.18. Eine Zahl $x \in \mathbb{R}$ heisst *Häufungswert* der reellen Folge (x_n) , falls diese eine gegen x konvergente Teilfolge besitzt. Ist (x_n) beschränkt, so ist die Menge der Häufungswerte nicht-leer wegen Theorem 14.17 und (offensichtlich) beschränkt. Man nennt die Zahlen

$$\begin{aligned} \limsup(x_n) &= \sup\{x \mid \text{Es existiert eine gegen } x \text{ konvergente Teilfolge von } (x_n)\} \\ \liminf(x_n) &= \inf\{x \mid \text{Es existiert eine gegen } x \text{ konvergente Teilfolge von } (x_n)\} \end{aligned} \quad (14.35)$$

den *limes superior* bzw. *limes inferior* von (x_n) . Für unbeschränkte Folgen verallgemeinert man entsprechend Definition 14.14.

Proposition 14.19. Sei $(x_n)_{n \in \mathbb{N}}$ eine reelle Folge.

- (i) Eine Zahl $x \in \mathbb{R}$ ist genau dann ein Häufungswert von (x_n) , wenn für jedes $\epsilon > 0$ unendlich viele Folgenglieder im Intervall $(x - \epsilon, x + \epsilon)$ liegen.
- (ii) Ist (x_n) (nach oben und unten) beschränkt, so existiert eine gegen $\limsup(x_n)$ und eine gegen $\liminf(x_n)$ konvergente Teilfolge, d.h.

$$\begin{aligned} \limsup(x_n) &= \max\{x \mid \text{Es existiert eine gegen } x \text{ konvergente Teilfolge von } (x_n)\} \\ \liminf(x_n) &= \min\{x \mid \text{Es existiert eine gegen } x \text{ konvergente Teilfolge von } (x_n)\} \end{aligned} \quad (14.36)$$

§ 14. FOLGEN REELLER ZAHLEN

(iii) Ist (x_n) (nach oben und unten) beschränkt, so gilt

$$\begin{aligned}\limsup(x_n) &= \min\{x \mid \forall \epsilon > 0 \text{ ist } x_n < x + \epsilon \text{ für fast alle } n.\} \\ \liminf(x_n) &= \max\{x \mid \forall \epsilon > 0 \text{ ist } x_n > x - \epsilon \text{ für fast alle } n.\}\end{aligned}\quad (14.37)$$

(iv) Im Allgemeinen (also auch ohne Beschränktheitsvoraussetzung) gilt:

$$\begin{aligned}\limsup(x_n) &= \inf\{x \mid x_n > x \text{ für höchstens endlich viele } n.\} \\ \liminf(x_n) &= \sup\{x \mid x_n < x \text{ für höchstens endlich viele } n.\}\end{aligned}\quad (14.38)$$

(v) Die Folge $(\bar{x}_n)$ mit $\bar{x}_n := \sup\{x_m \mid m \geq n\} \in \overline{\mathbb{R}} = \mathbb{R} \cup \{-\infty, +\infty\}$ ist monoton fallend, und

$$\limsup(x_n) = \lim_{n \rightarrow \infty} \bar{x}_n \in \overline{\mathbb{R}} \quad (14.39)$$

Analog: Die Folge $\underline{x}_n := \inf\{x_m \mid m \geq n\}$ ist monoton wachsend, und

$$\liminf(x_n) = \lim_{n \rightarrow \infty} \underline{x}_n \in \mathbb{R} \cup \{-\infty, +\infty\} \quad (14.40)$$

(vi) (x_n) konvergiert genau dann, wenn sie beschränkt ist und $\limsup(x_n) = \liminf(x_n)$.

Beweis. Wir fassen uns kurz.⁷⁹

(i) Folgt aus der Definition, haha.

(ii) Siehe Beweis von Theorem 14.17.

(iii) Folgt aus (i) und (ii).

(iv) Ditto. Ist etwa (x_n) nach oben unbeschränkt, so gilt $\limsup(x_n) = \inf \emptyset = +\infty$.

(v) Die Monotonie ist offensichtlich, der Grenzwert folgt aus (iv).

(vi) Klar. □

Zuletzt erreichen wir das Ziel des elegantesten und am meisten eingesetzten Konvergenzkriteriums.

Definition 14.20. Eine Folge reeller Zahlen (x_n) heisst Cauchy-Folge (CF), falls für jedes $\epsilon > 0$ ein $N_\epsilon \in \mathbb{N}$ existiert derart, dass

$$|x_n - x_m| < \epsilon \quad \text{falls } n \geq N_\epsilon \text{ und } m \geq N_\epsilon \quad (14.41)$$

Theorem 14.21. Eine Folge reeller Zahlen konvergiert dann und nur dann wenn sie eine Cauchy-Folge ist.

Bemerkungen. · Das Cauchy-Kriterium besagt intuitiv, dass Folgenglieder *voneinander* beliebig kleinen Abstand haben, wenn nur ihre Indizes gross genug sind. Dies ist ein enormer Vorteil gegenüber der ursprünglichen Definition 14.4, nach der man den Grenzwert *kennen* oder *raten* muss, bevor man testen kann, ob Folgenglieder ihm beliebig nahe kommen.

· Das Prinzip der monotonen Konvergenz 14.13 kommt auch ohne a priori Kenntnis des Grenzwert aus. Allerdings braucht es Prinzip die Anordnung von $\mathbb{R}$ in direkter

⁷⁹In der Praxis bietet die Identifikation von $\limsup/\liminf$ direkt aus der Definition selten grosse Schwierigkeiten.

Weise. Das Cauchy-Kriterium tut dies nur “indirekt”, nämlich über den Umweg des “Absolutbetrags”. Dies wird es uns ermöglichen, das Cauchy-Kriterium in einem sehr viel allgemeineren Kontext anzuwenden, in dem wir keine Anordnung, aber immer noch ein Analogon eines Absolutbetrags zur Verfügung haben.

· In der Definition einer Cauchy-Folge würde die Forderung $|x_n - x_{N_\epsilon}| < \epsilon \forall n \geq N_\epsilon$ ausreichen (Übungsaufgabe). Hingegen ist etwa schon die Forderung $|x_{n+1} - x_n| < \epsilon \forall n \geq N_\epsilon$ zu schwach.

Der Beweis des Cauchy-Kriteriums ist mit unseren Vorbereitungen relativ einfach. Wir verteilen ihn auf drei Lemmas.

Lemma 14.22. *Jede konvergente Folge ist eine Cauchy-Folge.*

Beweis. Sei $(x_n)_{n \in \mathbb{N}}$ konvergent mit Grenzwert x . Für $\epsilon > 0$ sei N_ϵ s.d. $|x - x_n| < \frac{\epsilon}{2} \forall n \geq N_\epsilon$. Dann gilt für alle $n \geq N_\epsilon$ und $m \geq N_\epsilon$,

$$|x_n - x_m| \leq |x_n - x| + |x - x_m| < \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon \quad (14.42)$$

□

Lemma 14.23. *Cauchy-Folgen sind beschränkt.*

Beweis. Sei $(x_n)_{n \in \mathbb{N}}$ eine Cauchy-Folge. Für $\epsilon = 1$ sei N_1 s.d. $|x_n - x_m| < 1 \forall n, m \geq N_1$. Insbesondere gilt mit $m = N_1$: $|x_n| = |x_{N_1} + x_n - x_{N_1}| \leq |x_{N_1}| + 1 \forall n \geq N_1$. Mit $B := \max\{|x_1|, |x_2|, \dots, |x_{N_1-1}|, |x_{N_1}| + 1\}$ gilt dann $|x_n| \leq B \forall n$. □

Lemma 14.24. *Jede Cauchy-Folge ist konvergent.*

Beweis. Nach Lemma 14.23 ist eine Cauchy-Folge (x_n) jedenfalls beschränkt. Nach 14.17 existiert eine konvergente Teilfolge (x_{n_k}) mit Grenzwert $x := \lim_{k \rightarrow \infty} x_{n_k}$. Wir behaupten, dass $x_n \xrightarrow{n \rightarrow \infty} x$.

Bew.: Für $\epsilon > 0$ sei N_ϵ derart, dass $|x_n - x_m| < \epsilon$ für alle $n \geq N_\epsilon$ und $m \geq N_\epsilon$. ((x_n) ist eine Cauchy-Folge.) Sodann sei K_ϵ derart, dass $n_{K_\epsilon} \geq N_\epsilon$ und $|x_{n_{K_\epsilon}} - x| < \epsilon$. (Konvergenz der Teilfolge). Zusammen folgt für alle $n \geq N_\epsilon$:

$$|x_n - x| \leq |x_n - x_{n_{K_\epsilon}}| + |x_{n_{K_\epsilon}} - x| < 2\epsilon \quad (14.43)$$

Daraus folgt die Behauptung. □

Mit 14.22 und 14.24 ist auch 14.21 bewiesen. □

Abschliessende Bemerkungen; Überabzählbarkeit von $\mathbb{R}$

· Das Cauchy-Kriterium ist ebenso wie das Intervallschachtelungsprinzip, unter der Voraussetzung des archimedischen Axiom, äquivalent zur Anordnungsvollständigkeit der reellen Zahlen.⁸⁰ Zur *Konstruktion* von $\mathbb{R}$ aus $\mathbb{Q}$ (siehe Theorem 6.15) sind Dedekindsche Schnitte jedoch ökonomischer (wenn nicht notwendiger Weise transparenter). Dies liegt daran, dass es zu jeder reellen Zahl x sehr viele verschiedene

⁸⁰Das werden wir hier nicht alles ausführen. Zur Implikation Cauchy-Kriterium $\Rightarrow$ Supremumseigenschaft siehe Beispiel 15.10

§ 14. FOLGEN REELLER ZAHLEN

(auch rationale) Intervallschachtelungen mit Durchschnitt x gibt, und ebenso sehr viele (rationale) Cauchy-Folgen mit Grenzwert x . Die reellen Zahlen müssten also als Quotient der Menge der (rationalen) Intervallschachtelungen/Cauchy-Folgen bezüglich einer geeigneten Äquivalenzrelation definiert werden.

· Während die Definitionen und Rechenregeln zur Konvergenz auch für Folgen rationaler Zahlen gegolten hätten, hängen die Konvergenzkriterien 14.13 (monotone Konvergenz), 14.21 (Cauchy-Kriterium) sowie die Ergebnisse 14.16 (Intervallschachtelungen) sowie 14.17 (Satz von Bolzano-Weierstrass) wesentlich von der Vollständigkeit der reellen Zahlen ab. Der Preis dafür, und ein Mass für die “Irrationalität” der reellen Zahlen, ist ihre Überabzählbarkeit.

Erinnerung. (Siehe Def. 4.10) Eine Menge X ist abzählbar, wenn eine injektive Abbildung von X nach $\mathbb{N}$ existiert. Eine solche Menge X ist endlich, wenn das Bild einer solchen Abbildung beschränkt ist, andernfalls abzählbar unendlich. Äquivalent dazu ist X genau dann abzählbar unendlich, wenn eine surjektive Abbildung $\mathbb{N} \rightarrow X$ existiert, mit anderen Worten ist dies genau eine Folge $(a_n) \subset X$ mit der Eigenschaft, dass jedes Element $x \in X$ von (mindestens) einem a_n getroffen wird. Durch Übergang zu einer Teilfolge erhält man daraus eine bijektive Abbildung zwischen $\mathbb{N}$ und X .

· Die Menge der rationalen Zahlen ist abzählbar. Eine Abzählung entsteht zum Beispiel, indem man (nach der Null) sukzessive für jedes $N \in \mathbb{N}$ die endlich vielen rationalen Zahlen aufreihet, deren (gekürzte) Zähler und Nenner beide nicht grösser als N sind.

· Man zeige: Eine Abzählung von $\mathbb{Q}$ hat jede reelle Zahl als Häufungswert. Man sagt, die “rationalen Zahlen liegen dicht in $\mathbb{R}$ ”.

Theorem 14.25. *Die Menge der reellen Zahlen ist nicht abzählbar.*

Beweis. Wir zeigen, dass für jede Folge (x_n) paarweise verschiedener reeller Zahlen mindestens eine reelle Zahl h existiert mit $h \neq x_n \forall n \in \mathbb{N}$. (Insbesondere kann also keine Abzählung *aller* reeller Zahlen existieren.)

Bew.: · Sei $a_1 = \min\{x_1, x_2\}$, $b_1 = \max\{x_1, x_2\}$ und $I_1 = (a_1, b_1)$ (das offene Intervall!) Wir setzen $k_1 = 1$ oder $k_1 = 2$ so, dass $a_1 = x_{k_1}$, und entsprechend l_1 so, dass $b_1 = x_{l_1}$. Wegen $a_1 < b_1$ ist I_1 nicht leer und enthält unendlich viele reelle Zahlen. Falls nur endlich viele Glieder von (x_n) in I_1 liegen, so existiert bereits ein $h \in I_1$, welches nicht in (x_n) vorkommt.

· Wir nehmen daher an, dass unendlich viele Glieder von (x_n) in I_1 liegen, und bezeichnen die ersten beiden mit y_1 und z_1 . Sei $a_2 = \min\{y_1, z_1\}$ und $b_2 = \max\{y_1, z_1\}$. Dann gilt $a_1 < a_2 < b_2 < b_1$ und wir schreiben $I_2 = (a_2, b_2) \subsetneq I_1$. Es existiert (genau) ein k_2 mit $a_2 = x_{k_2}$ und (genau) ein l_2 mit $b_2 = x_{l_2}$.

· Wir wiederholen das Verfahren und erhalten entweder ein h welches nicht von (x_n) getroffen wird oder aber eine streng monoton wachsende Teilfolge $(a_m) = (x_{k_m})$ und eine streng monoton fallende Teilfolge $(b_m) = (x_{l_m})$. (Die Monotonie von (k_m) und (l_m) folgt aus der Wahl von y_m, z_m als die ersten Glieder von (x_n) in I_m und der Schachtelung der Intervalle.)

· Wegen 14.13 existiert $a = \lim_{m \rightarrow \infty} a_m$ und $b = \lim_{m \rightarrow \infty} b_m$. Es gilt $a \leq b$ und es existiert ein $h \in [a, b]$.

- Wegen der strengen Monotonie von (a_m) und (b_m) liegt h in keiner dieser beiden Folgen. Wir behaupten, h liegt auch nicht in der ursprünglichen Folge (x_n) .
- Denn angenommen $h = x_{n_*}$, dann treten nur endlich viele $a_m = x_{k_m}$ vor h in der Folge (x_n) auf, d.h. es gibt ein letztes d mit $k_d < n_*$ aber $k_m > n_*$ $\forall m > d$.
- Insbesondere gilt $k_{d+1} > n_*$, d.h. a_{d+1} tritt in der Folge (x_n) nach $x_{n_*} = h$ auf. Wegen $h \in (a_m, b_m)$ $\forall m$ steht dies aber im Widerspruch dazu, dass a_{d+1} und b_{d+1} die ersten Glieder von (x_n) im Intervall (a_d, b_d) sind. $\square$
- Es gibt einfachere Beweise mittels Entwicklung reeller Zahlen in Brüche mit einer festen Basis (z.B. Dezimal- oder Binärbrüche).

§ 15 Metrische Räume

In § 8 (siehe auch die Einführung zu Kapitel 2) hatten wir als Grundthese der wissenschaftlichen Naturbeschreibung den *approximativen* Vergleich realer Gegebenheiten mit *idealisierten* mathematischen Strukturen vorgeschlagen, und die operationelle Begründung des affinen euklidischen Raums wenigstens skizziert. In diesem § werden die *Abstandsbegriffe* (nach-)geliefert, die es erlauben sollten, die genannten Approximationen (a.k.a., *Messunsicherheiten*) zu *quantifizieren* und (so die Hoffnung) ständig zu *verbessern*. Dies führt im mathematischen Bild auf die angekündigte Übertragung des Vollständigkeitsbegriffs auf Situationen, in denen eine Anordnung entweder nicht vorhanden oder sogar gänzlich unmöglich ist. Wenn (sich herausstellt, dass) die fragliche physikalische Grösse eine (affin-)lineare Struktur besitzt, sollten die Abstandsbegriffe mit den Vektorraumoperationen verträglich sein. Hier ist die wesentliche Erkenntnis, dass für endlich-dimensionale (reelle) Vektorräume die resultierenden Konvergenz- und Vollständigkeitsbegriffe *eindeutig* sind (im unendlichen-dimensionalen Fall allerdings nicht, s. HöMa 2 u. 3).

Definition 15.1. Ein *metrischer Raum* (sc. über $\mathbb{R}$) ist eine Menge X zusammen mit einer *Abstandsfunktion* oder *Metrik*⁸¹ genannten Abbildung

$$d_X = d : X \times X \rightarrow \mathbb{R}_{\geq 0} = \{x \in \mathbb{R} \mid x \geq 0\} \quad (15.1)$$

mit den Eigenschaften

$$\begin{aligned} \text{Symmetrie: } d(x_1, x_2) &= d(x_2, x_1) \quad \forall x_1, x_2 \in X \\ \text{Positivität: } d(x_1, x_2) &= 0 \Leftrightarrow x_1 = x_2 \end{aligned} \quad (15.2)$$

$$\text{Dreiecksungleichung: } d(x_1, x_2) \leq d(x_1, x_3) + d(x_3, x_2) \quad \forall x_1, x_2, x_3 \in X$$

Definition 15.2. Eine *Norm* auf einem reellen Vektorraum V ist eine Funktion

$$\|\cdot\| : V \rightarrow \mathbb{R}_{\geq 0} \quad (15.3)$$

mit den Eigenschaften

$$\begin{aligned} \text{(absolute) Homogenität: } \|\alpha v\| &= |\alpha| \cdot \|v\| \quad \forall \alpha \in \mathbb{R}, v \in V \\ \text{Positivität: } \|v\| &> 0 \quad \forall v \neq 0 \\ \text{Dreiecksungleichung: } \|v_1 + v_2\| &\leq \|v_1\| + \|v_2\| \quad \forall v_1, v_2 \in V \end{aligned} \quad (15.4)$$

⁸¹In der Physik ist der Begriff “Metrik” für das vorbehalten, was in der Mathematik später als “metrischer Tensor” bezeichnet wird.

§ 15. METRISCHE RÄUME

Ein solches Paar $(V, \|\cdot\|)$ heisst *normierter Vektorraum* (manchmal auch nur *normierter Raum*).

Bemerkung. Prinzipiell liesse sich als Wertebereich für Abstandsfunktionen und Normen jeder beliebige angeordnete Körper ins Auge fassen. Eine zufrieden stellende Konvergenztheorie für Folgen (und zwar insbesondere Argumente, die über “ $\epsilon = \frac{1}{n}$ ” laufen) erfordert allerdings die Gültigkeit des archimedischen Axioms, welches gemäss Theorem 6.15 (iii) nur in Unterkörpern von $\mathbb{R}$ (wie etwa $\mathbb{Q}$) erfüllt ist. Als Skalarbereich (a.k.a., “Grundkörper”) für normierte Vektorräume kommt jeder Körper in Frage, auf dem ein Absolutbetrag (mit Werten in $\mathbb{R}$ und den gleichen Eigenschaften wie insbesondere auf $\mathbb{C}$) definiert werden kann und der dann in Gleichung (15.4) an die Stelle von $|\cdot|$ tritt.

Beispiel 15.3. 1. Der n -dimensionale reelle Raum aus Definition 8.1,

$$\mathbb{R}^n = \left\{ v = (v^1, \dots, v^n)^T = \begin{pmatrix} v^1 \\ \vdots \\ v^n \end{pmatrix} \mid v^i \in \mathbb{R} \right\} \quad (15.5)$$

ausgerüstet mit der “euklidischen Standardnorm” (s. Def. 8.6)

$$\|v\|_2 := \left(\sum_{i=1}^n (v^i)^2 \right)^{1/2} \quad (15.6)$$

ist ein normierter Vektorraum. (Zur Erinnerung: Die Dreiecksungleichung folgt aus der Cauchy-Schwarz-Ungleichung Prop. 8.7 für das zugehörige euklidische innere Produkt Def. 8.4.)

Obwohl geometrisch anschaulich und wohlvertraut ist die euklidische Norm (u.a. wegen der Wurzel!) nicht für alle Zwecke am nützlichsten. Hier sind einige andere.

Beispiele 15.4. Für jedes Tupel $\lambda = (\lambda_i)_{i=1, \dots, n}$ positiver reeller Zahlen ist

$$\|v\|_{2,\lambda} := \left(\sum_{i=1}^n \lambda_i (v^i)^2 \right)^{1/2} \quad (15.7)$$

ebenfalls eine Norm auf $\mathbb{R}^n$. (Beweis durch Rückführung auf (15.6) mittels einer geeigneten Koordinatentransformation.)

Die Funktion $\|\cdot\|_\infty$, definiert durch

$$\|v\|_\infty := \max\{|v^i| \mid i = 1, \dots, n\} \quad (15.8)$$

ist eine Norm auf $\mathbb{R}^n$. Homogenität und Positivität sind klar, die Dreiecksungleichung folgt aus 6.9 via

$$\|v_1 + v_2\|_\infty = \max\{|v_1^i + v_2^i|\} \leq \max\{|v_1^i|\} + \max\{|v_2^i|\} = \|v_1\|_\infty + \|v_2\|_\infty \quad (15.9)$$

Ausführlicher: Es gilt $\max\{|v_1^i + v_2^i|\} = |v_1^{i_*} + v_2^{i_*}|$ für ein gewisses $i_* \in \{1, \dots, n\}$. Für dieses i_* gilt zunächst $|v_1^{i_*} + v_2^{i_*}| \leq |v_1^{i_*}| + |v_2^{i_*}|$ wegen der gewöhnlichen Dreiecksungleichung, und dann weiters $|v_1^{i_*}| \leq |v_1^i| \ \forall i$, d.h. $|v_1^{i_*}| \leq \max\{|v_1^i|\}$, ebenso $|v_2^{i_*}| \leq \max\{|v_2^i|\}$, zusammen dann (15.9).

· Die Abbildung $\|\cdot\|_1$, definiert durch

$$\|v\|_1 := \sum_{i=1}^n |v^i| \quad (15.10)$$

ist eine Norm auf $\mathbb{R}^n$. Homogenität und Positivität sind wieder klar, die Dreiecksungleichung folgt wie oben aus (6.9):

$$\|v_1 + v_2\|_1 = \sum_{i=1}^n |v_1^i + v_2^i| \leq \sum_{i=1}^n (|v_1^i| + |v_2^i|) = \sum_{i=1}^n |v_1^i| + \sum_{i=1}^n |v_2^i| = \|v_1\|_1 + \|v_2\|_1 \quad (15.11)$$

Beispiel 15.5. Für jede rationale Zahl $p \in \mathbb{Q}$, $p > 1$ ist die Abbildung $\|\cdot\|_p$, definiert durch

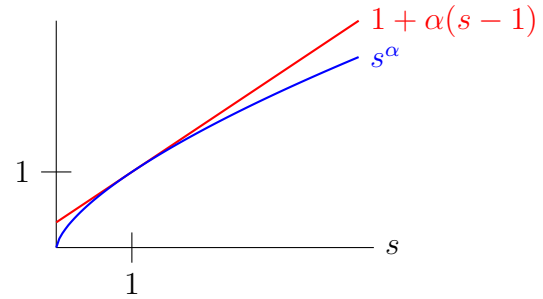
$$\|v\|_p = \left(\sum_{i=1}^n |v^i|^p \right)^{1/p} \quad (15.12)$$

eine Norm auf $\mathbb{R}^n$ (die sog. p -Norm). Homogenität und Positivität sind wieder klar. Herleitung der Dreiecksungleichung:

1. Schritt: Für alle $s > 0$, $\alpha \in \mathbb{Q}$ mit $\alpha < 1$ gilt (vgl. Lemma 6.3)

$$s^\alpha \leq 1 + \alpha(s - 1) \quad (15.13)$$

Bew. Die Funktion $f : (0, \infty) \rightarrow \mathbb{R}_{>0}$ mit $f(s) := s^\alpha - (1 + \alpha(s - 1))$ ist zweimal differenzierbar mit $f(1) = 0$, $f'(1) = 0$ und $f''(s) = \alpha(\alpha - 1)s^{\alpha-2} < 0 \ \forall s \in (0, \infty)$. Aus elementarer Differentialrechnung folgt $f(s) \leq 0 \ \forall s \in (0, \infty)$.



2. Schritt: Für alle $x, y \geq 0$, α wie oben gilt mit $\beta = 1 - \alpha$ die *Ungleichung zwischen dem (gewichteten) arithmetischen und geometrischen Mittel*

$$x^\alpha y^\beta \leq \alpha x + \beta y \quad (15.14)$$

Bew.: Für $y = 0$ ist die Sache klar. Andernfalls erhalten wir durch Einsetzen von $s = x/y$ in (15.13)

$$\left(\frac{x}{y}\right)^\alpha \leq 1 + \alpha\left(\frac{x}{y} - 1\right) \quad (15.15)$$

und daraus durch Multiplikation mit y unmittelbar (15.14).

3. Schritt: Für alle $v, w \in \mathbb{R}^n$, $p, q > 1$ mit $\frac{1}{p} + \frac{1}{q} = 1$ gilt die *Höldersche Ungleichung*

$$\left| \sum_{i=1}^n v^i w^i \right| \leq \|v\|_p \cdot \|w\|_q \quad (15.16)$$

Bew.: Für $v = 0$ oder $w = 0$ ist die Sache klar. Andernfalls erhalten wir mit $\alpha = 1/p$, $\beta = 1/q$ durch Einsetzen von $x = (|v^i|/\|v\|_p)^p$, $y = (|w^i|/\|w\|_q)^q$ in (15.14)

$$\frac{|v^i|}{\|v\|_p} \cdot \frac{|w^i|}{\|w\|_q} \leq \frac{1}{p} \frac{|v^i|^p}{(\|v\|_p)^p} + \frac{1}{q} \frac{|w^i|^q}{(\|w\|_q)^q} \quad (15.17)$$

§ 15. METRISCHE RÄUME

für alle i , und daraus durch Summation über i

$$\frac{1}{\|v\|_p \cdot \|w\|_q} \left| \sum_{i=1}^n v^i w^i \right| \leq \sum_{i=1}^n \frac{|v^i|}{\|v\|_p} \cdot \frac{|w^i|}{\|w\|_q} \leq \frac{1}{p} + \frac{1}{q} = 1 \quad (15.18)$$

Daraus folgt (15.16).

4. Schritt: Wir rechnen:

$$\begin{aligned} (\|v + w\|_p)^p &= \sum_{i=1}^n |v^i + w^i|^{p-1} |v^i + w^i| \\ \text{(gewöhnliche Dreiecksungl.)} \quad &\leq \sum_{i=1}^n |v^i + w^i|^{p-1} |v^i| + \sum_{i=1}^n |v^i + w^i|^{p-1} |w^i| \\ \text{(wegen Hölder (15.16))} \quad &\leq \left(\sum_{i=1}^n |v^i + w^i|^{(p-1)q} \right)^{1/q} (\|v\|_p + \|w\|_p) \\ \text{(wegen } pq - q = p) \quad &= (\|v + w\|_p)^{p/q} \cdot (\|v\|_p + \|w\|_p) \end{aligned} \quad (15.19)$$

Daraus folgt wegen $p - p/q = 1$ die als *Minkowski-Ungleichung*

$$\|v + w\|_p \leq \|v\|_p + \|w\|_q \quad (15.20)$$

bekannte Dreiecksungleichung für die p -Norm.

Proposition 15.6. Sei (X, τ, δ) ein affiner Raum⁸² über dem reellen Vektorraum V .

(i) Ist $\|\cdot\|$ eine Norm auf V , so ist

$$\begin{aligned} d_{\|\cdot\|} : X \times X &\rightarrow \mathbb{R}_{\geq 0} \\ d_{\|\cdot\|}(x_1, x_2) &:= \|\delta(x_1, x_2)\| \end{aligned} \quad (15.21)$$

eine Abstandsfunktion auf X . Sie erfüllt die weiteren Eigenschaften

$$\begin{aligned} \text{Translationsinvarianz:} \quad d_{\|\cdot\|}(\tau_v(x_1), \tau_v(x_2)) &= d_{\|\cdot\|}(x_1, x_2) \\ \text{Homogenität:} \quad d_{\|\cdot\|}(x_1, \tau_{\alpha\delta(x_1, x_2)}(x_1)) &= |\alpha| \cdot d_{\|\cdot\|}(x_1, x_2) \end{aligned} \quad (15.22)$$

(ii) Ist umgekehrt d_X eine Abstandsfunktion auf X mit diesen Eigenschaften (15.22), $x_0 \in X$, so ist $\|v\|_d := d_X(x_0, \tau_v(x_0))$ eine (von x_0 unabhängige) Norm auf V .

(iii) Diese Aussagen gelten insbesondere für $V = X$ mit $\delta(v_1, v_2) = v_2 - v_1$ als affinen Raum über sich selbst.

Beweis. Umschreiben der drei Eigenschaften aus Def. 15.2 auf die in Def. 15.1 und Gl. (15.22), und umgekehrt, unter Zuhilfenahme von Def. 8.2 und Lemma 8.3. $\square$

⁸²Erinnerung: $\tau : V \times X \rightarrow X$ sind die Translationen, $\delta : X \times X \rightarrow V$ der Abstandsvektor. Es gilt $\tau_{v_1} \circ \tau_{v_2} = \tau_{v_1+v_2}$, $\delta(x_0, \tau_v(x_0)) = v$, $\tau_{\delta(x_1, x_2)}(x_1) = x_2$, $\delta(x_1, x_2) + \delta(x_2, x_3) = \delta(x_1, x_3)$.

Beispiel 15.7. · Auf der Einheitssphäre

$$S^2 = \{x \in \mathbb{R}^3 \mid \|x\|_2 = 1\} \subset \mathbb{R}^3 \quad (15.23)$$

im drei-dimensionalen euklidischen Raum ist $d_{S^2}(x, y) := \arccos \langle x, y \rangle \in [0, \pi]$, der *orthodromische Abstand*, definiert als die Länge des kürzeren Grosskreisabschnittes zwischen x und y , erhältlich als Schnitt von S^2 mit der durch $x \neq y$ und dem Ursprung aufgespannten Ebene (auch bekannt als geodätischer Abstand). Beachte hierzu, dass wegen der C.-S.U. (8.7) $\forall x, y \in S^2 \langle x, y \rangle \in [-1, 1]$, mit eindeutigem Zwischenwinkel in $[0, \pi]$. Wir zeigen hier nicht die Dreiecksungleichung, bemerken aber, dass in diesem Beispiel die Bildmenge von d_{S^2} beschränkt ist.

· Auf jede Menge X ist

$$d_X(x, y) = \begin{cases} 0 & \text{falls } x = y \\ 1 & \text{sonst} \end{cases} \quad (15.24)$$

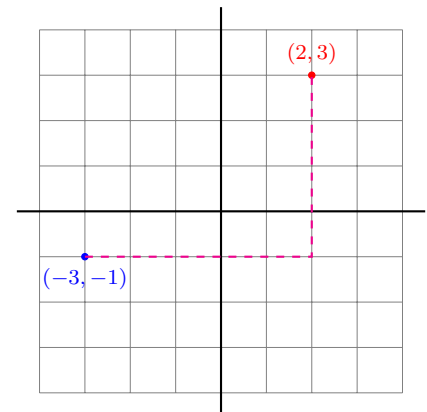
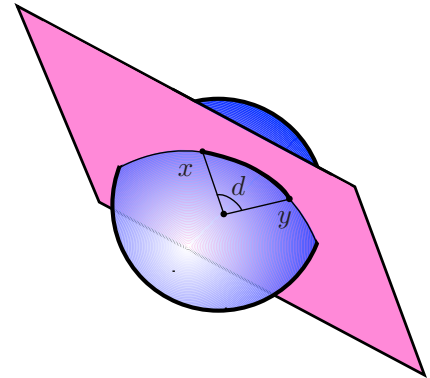
eine Abstandsfunktion und damit (X, d_X) ein metrischer Raum.

· Ein praktisches Beispiel ist $\mathcal{B} = (\text{Menge der Bahnhöfe der Deutschen Bahn})$ mit

$$d_{\mathcal{B}}(A, B) = \text{Dauer der kürzesten Verbindung von } A \text{ nach } B \quad (15.25)$$

· Die Einschränkung der durch die Eins-Norm (15.10) auf $\mathbb{R}^2$ induzierten Metrik auf ein diskretes Gitter definiert die sog. “Manhattan⁸³-Metrik”, im Beispiel

$$d((-3, -1), (2, 3)) = |2 - (-3)| + |3 - (-1)| = 9 \quad (15.26)$$



Wir können nun daran gehen, den Vollständigkeitsbegriff auf metrische Räume (über $\mathbb{R}$) und damit via Prop. 15.6 insbesondere auf normierte Vektorräume auszudehnen. Hierzu erinnern wir daran, dass wir (14.1) Folgen (und Teilfolgen) mit Werten in allgemeinen Mengen X eingeführt hatten, und für $X = \mathbb{R}$ an die Definition einer Cauchy-Folge in 14.20. Die wesentliche Beobachtung ist, dass im Cauchy-Kriterium 14.21 die Anordnung von $\mathbb{R}$ nur mittels des Absolutbetrags $|\cdot| : \mathbb{R} \rightarrow \mathbb{R}_{\geq 0}$ eingeht und in die meisten Konvergenzargumente nur dessen metrische Eigenschaften (Dreiecksungleichung und Positivität).

Definition 15.8. (i) Eine Folge $(x_k)_{k \in \mathbb{N}}$ in einem metrischen Raum (X, d) heisst *konvergent*, falls ein $x \in X$ existiert⁸⁴ mit der Eigenschaft, dass $\forall \epsilon > 0$ ein $K_\epsilon \in \mathbb{N}$

⁸³Laut Wikipedia auch “Mannheimer” genannt.

⁸⁴Elemente von X werden normalerweise als Punkte adressiert, Limites von Folgen aber oft wieder als “Grenzwerte”.

§ 15. METRISCHE RÄUME

existiert so, dass

$$d(x, x_k) < \epsilon \quad \forall k \geq K_\epsilon \quad (15.27)$$

(ii) Eine Folge (x_k) heisst *Cauchy-Folge*, falls $\forall \epsilon > 0$ ein K_ϵ existiert so, dass

$$d(x_k, x_l) < \epsilon \quad \text{falls } k \geq K_\epsilon \text{ und } l \geq K_\epsilon \quad (15.28)$$

(iii) Ein metrischer Raum heisst *Cauchy-vollständig* (oder nur *vollständig*), falls jede Cauchy-Folge in X konvergiert.

Definition 15.9. Ein normierter Vektorraum $(V, \|\cdot\|)$ heisst *vollständig*, wenn der gemäss 15.6 (iii) assoziierte metrische Raum vollständig ist, d.h. im Klartext:

$$\begin{aligned} \forall (v_k) \subset V : & \left((\forall \epsilon > 0 \exists K_\epsilon : k, l \geq K_\epsilon \Rightarrow \|v_k - v_l\| < \epsilon) \right. \\ & \left. \Rightarrow (\exists v \in X : \forall \epsilon > 0 \exists K_\epsilon : k \geq K_\epsilon \Rightarrow \|v_k - v\| < \epsilon) \right) \end{aligned} \quad (15.29)$$

Ein vollständiger normierter Vektorraum heisst auch *Banach-Raum*.

Bemerkung. Wie bei $\mathbb{R}$ zeigt man leicht, dass Grenzwerte eindeutig sind (dies basiert auf $d(x, y) < \epsilon \quad \forall \epsilon > 0 \Rightarrow x = y$, vgl. S. 95), und dass jede konvergente Folge eine Cauchy-Folge ist (Dreiecksungleichung!). Der wesentliche Punkt ist also Cauchy-Folge $\Rightarrow$ Konvergenz.

· Man beachte auch wieder, dass eine Folge genau dann gegen $x \in X$ konvergiert wenn $d(x, x_k)$ eine Nullfolge ist, s. 14.6. (Der Versuch einer entsprechende Charakterisierung von Cauchy-Folgen in X über ihr Bild unter d wäre aber offenbar unsinnig.)

· Für eine konvergente Folge $(x_n) \subset X$ in einem metrischen Raum schreiben wir ebenfalls $\lim_{n \rightarrow \infty} x_n \in X$ für ihren Grenzwert.

Beispiel 15.10. Die Menge der reellen Zahlen $\mathbb{R}$, ausgestattet mit der Abstandsfunktion $d_{\mathbb{R}}(x, y) = |x - y|$, ist ein vollständiger metrischer Raum im Sinne von 15.8 (siehe Theorem 14.21). Um mit der Anordnungsvollständigkeit aus Def. 6.14 zu vergleichen, und damit eine weitere der auf S. 102 angekündigten Implikationen zu zeigen, bemerken wir Folgendes: Ist K ein angeordneter archimedischer Körper, so können wir den Absolutbetrag $|\cdot|_K$, definiert über (6.10), nach Thm. 6.15 auffassen als Abbildung $K \rightarrow \mathbb{R}_{\geq 0}$, und man prüft sofort, dass durch $d_K(x, y) = |x - y|_K$ eine Abstandsfunktion im Sinne von Def. 15.1 erklärt wird. Es gilt dann:

Beh.: Ist K mit dieser Abstandsfunktion ein vollständiger metrischer Raum, so ist K anordnungsvollständig, d.h. jede nicht-leere nach oben beschränkte Menge besitzt eine kleinste obere Schranke. M.a.W. ist $K = \mathbb{R}$.

Bew.: (Skizze) Sei $T \subset K$ nicht leer und $T \leq B$ für ein $B \in K$. Wähle $x \in T$ und setze $[a_1, b_1] = [x - 1, B]$ (so dass garantiert $(x - 1, B) \neq \emptyset$). Definiere dann rekursiv für $k = 1, 2, \dots$

$$\begin{aligned} m_k &:= \frac{a_k + b_k}{2} \\ [a_{k+1}, b_{k+1}] &:= \begin{cases} [a_k, m_k] & \text{falls } [m_k, b_k] \cap T = \emptyset \\ [m_k, b_k] & \text{falls } [m_k, b_k] \cap T \neq \emptyset \end{cases} \end{aligned} \quad (15.30)$$

Die Folge $([a_k, b_k])_{k \in \mathbb{N}}$ ist dann eine Intervallschachtelung in K (!) (Hier ist es wichtig, dass $|b_k - a_k|_K = 2^{1-k}|b_1 - a_1|_K$ wegen 6.11 bereits in K eine Nullfolge ist, da K archimedisch ist.), und man zeigt ohne grosse Mühe, dass (a_k) , (b_k) Cauchy-Folgen sind, deren gemeinsamer Grenzwert $s = \lim a_k = \lim b_k$ eine kleinste obere Schranke von T ist. (Alle b_k sind obere Schranken von T , und für jedes $\epsilon > 0$ ist $b_k - \epsilon < a_k$ für k gross genug keine obere Schranke, da $[a_k, b_k] \cap T \neq \emptyset \forall k$.)

Beispiel 15.11. Die Beispiele (15.24) und (15.25) aus 15.7 sind vollständige metrische Räume (Übungsaufgabe). Ebenso ist (15.23) ein vollständiger metrischer Raum, was wir aber noch nicht ganz beweisen können.

- Der Körper der rationalen Zahlen ist mit dem Absolutbetrag ausgestattet als metrischer Raum (über $\mathbb{R}$) aufgefasst nicht vollständig.
- Die natürlichen Zahlen $\mathbb{N}$ mit der Abstandsfunktion $d(n, m) = |\frac{1}{n} - \frac{1}{m}|$ sind ein unvollständiger metrischer Raum: Die Folge $x_n = n$ ist eine Cauchy-Folge ohne Grenzwert in $\mathbb{N}$.
- Andererseits ist $\mathbb{N}$ mit der Abstandsfunktion $d(n, m) = |n - m|$ ein vollständiger metrischer Raum: Ab $\epsilon = 1$ ist jede Cauchy-Folge konstant, also insbesondere konvergent.

Proposition 15.12. Der Vektorraum $\mathbb{R}^n$ mit der Norm $\|\cdot\|_\infty$ aus (15.8) ist vollständig.

Beweis. Sei $(x_k)_{k \in \mathbb{N}} \subset \mathbb{R}^n$ eine Cauchy-Folge.⁸⁵ Für beliebiges $\epsilon > 0$ sei K_ϵ so, dass $\|x_k - x_l\|_\infty < \epsilon$ falls $k, l \geq K_\epsilon$. Dann folgt aus der für alle $i = 1, \dots, n$ und alle $y \in \mathbb{R}^n$ gültigen Abschätzung

$$|y^i| \leq \max\{|y^i| \mid i = 1, \dots, n\} = \|y\|_\infty \quad (15.31)$$

dass $|x_k^i - x_l^i| \leq \|x_k - x_l\|_\infty < \epsilon \forall k, l \geq K_\epsilon$. Für jedes i ist also die Folge der Komponenten $(x_k^i)_{k \in \mathbb{N}}$ eine reelle Cauchy-Folge. Wir setzen $x^i = \lim_{k \rightarrow \infty} x_k^i$ (gemäss Thm. 14.21) und behaupten, dass $\lim_{k \rightarrow \infty} x_k = x := (x^1, \dots, x^n)^T$, jetzt im Sinne von Def. 15.8. Für $\epsilon > 0$ seien dazu für $i = 1, \dots, n$ K_ϵ^i so, dass $|x_k^i - x^i| < \epsilon$ falls $k \geq K_\epsilon^i$. Mit $K_\epsilon := \max\{K_\epsilon^i \mid i = 1, \dots, n\}$ gilt dann $\|x - x_k\|_\infty = \max\{|x^i - x_k^i|\} < \epsilon$ für alle $k \geq K_\epsilon$. $\square$

In ganz ähnlicher Weise kann man zeigen, dass auch $(\mathbb{R}^n, \|\cdot\|_2)$ vollständig ist, und unsere eingangs gestellte Aufgabe ist damit weitgehend erfüllt. Vor weiteren Beispielen wollen wir aber noch im Rahmen der allgemeinen Theorie der metrischen und normierten Räume der Frage nach der Eindeutigkeit der Konvergenz- und Vollständigkeitsbegriffe nachgehen. Ohne weiteres lässt sich hier zwar nur wenig sagen, wie das Beispiel von $\mathbb{N}$ oben illustriert. Die Forderung nach Verträglichkeit mit algebraischen Strukturen aber schränkt die Möglichkeiten wieder stark ein, und für endlich-dimensionale reelle Vektorräume ist das Resultat tatsächlich eindeutig.

Mit Blick auf S. 95, dass wir zur Konvergenz einer Folge das “ $\epsilon > 0$ nur bis auf einen konstanten Faktor schlagen müssen”, erklären wir:

⁸⁵Wir bezeichnen Elemente des $\mathbb{R}^n$ ab jetzt manchmal wieder mit $x = (x^i)_{i=1, \dots, n}$ statt v , unter anderem um anzudeuten, dass wir eher an einen (affinen, euklidischen) physikalischen Raum denken als einen abstrakten Vektorraum.

§ 15. METRISCHE RÄUME

Definition 15.13. Sei V ein reeller Vektorraum. Zwei Normen $\|\cdot\|_\alpha$ und $\|\cdot\|_\beta$ auf V heißen *äquivalent*, falls reelle Konstanten $c > 0$ und $C > 0$ existieren so dass für alle $v \in V$

$$c \cdot \|v\|_\alpha \leq \|v\|_\beta \leq C \cdot \|v\|_\alpha \quad (15.32)$$

Lemma 15.14. Äquivalenz von Normen ist eine Äquivalenzrelation auf der Menge der Normen auf V .

Beweis. Reflexivität und Transitivität sind offensichtlich. Symmetrie folgt durch Umschreiben der Ungleichungen (15.32) auf $\frac{1}{C} \cdot \|v\|_\beta \leq \|v\|_\alpha \leq \frac{1}{c} \cdot \|v\|_\beta$. $\square$

Es gilt dann:

Lemma 15.15. (i) Eine Folge $(x_n) \subset V$ konvergiert genau dann bezüglich $\|\cdot\|_\alpha$ (gemeint ist hier: im metrischen Raum (V, d_α) , wo d_α die zu $\|\cdot\|_\alpha$ gehörige Abstandsfunktion ist, man sagt auch “konvergiert in der Norm $\|\cdot\|_\alpha$ ”), wenn sie bezüglich $\|\cdot\|_\beta$ konvergiert. Der Grenzwert ist der gleiche.

(ii) Eine Folge $(x_n) \subset V$ ist genau dann eine Cauchy-Folge bezüglich $\|\cdot\|_\alpha$ wenn sie eine Cauchy-Folge bezüglich $\|\cdot\|_\beta$ ist.

(iii) $(V, \|\cdot\|_\alpha)$ ist genau dann vollständig, wenn $(V, \|\cdot\|_\beta)$ vollständig ist.

wie man leicht verifiziert. $\square$

Für ein Andermal: Für metrische Räume führt das Ersetzen der Normen in (15.32) durch die Abstandsfunktionen auf den Begriff der “strengen Äquivalenz”, der zwar sinnvoll ist, i.A. aber stärker als der Vergleich über Konvergenz von Folgen bzw. der induzierten Topologien, bei dem die konstanten c, C noch von Punkt zu Punkt variieren können. Gilt in (15.32) nur die zweite Ungleichung, so heisst $\|\cdot\|_\alpha$ *stärker* als $\|\cdot\|_\beta$. Konvergenz in $\|\cdot\|_\alpha$ impliziert dann Konvergenz in $\|\cdot\|_\beta$, aber nicht notwendig umgekehrt.

Beispiel 15.16. Auf $\mathbb{R}^n$ sind die ∞ - und 2-Norm äquivalent:

- Aus $|x^i|^2 \leq \sum_{j=1}^n |x^j|^2 \ \forall i$ folgt $|x^i| \leq \|x\|_2 \ \forall i$ und daher $\|x\|_\infty \leq \|x\|_2$.
- Andererseits gilt

$$\sum_{i=1}^n |x^i|^2 \leq n \cdot \max\{|x^i|^2\} = n \|x\|_\infty^2 \quad (15.33)$$

also $\|x\|_2 \leq \sqrt{n} \|x\|_\infty$. Es gilt also (15.32) mit $\alpha = \infty$, $\beta = 2$, $c = 1$ und $C = \sqrt{n}$.

- Mit Lemma 15.15 folgt insbesondere, dass auch der euklidische Raum mit Abstandsfunktion (8.9) vollständig ist. Dies ist kein Zufall:

Theorem 15.17. Je zwei Normen auf einem endlich-dimensionalen (reellen oder komplexen) Vektorraum sind äquivalent.

Zur Vorbereitung des Beweises etwas Kontext.

Topologische Grundbegriffe

Die grundlegende Beobachtung ist die folgende: Sei (X, d) ein metrischer Raum, und $T \subset X$ eine Teilmenge. Dann erfüllt die Einschränkung

$$d_T = d|_{T \times T} : T \times T \rightarrow \mathbb{R} \quad (15.34)$$

trivialerweise wieder alle Eigenschaften einer Abstandsfunktion, d.h. $(T, d|_{T \times T})$ ist in natürlicher Weise ein metrischer Raum. Beim Vergleich zwischen Konvergenz in T und Konvergenz in X trifft man dann auf die folgende Unterscheidung.

Definition 15.18. (i) Eine Teilmenge A eines metrischen Raumes heisst *abgeschlossen*⁸⁶, falls für jede Folge $(x_k)_{k \in \mathbb{N}}$ in A , welche als Folge in X konvergiert, der Grenzwert $x = \lim x_k$ in A liegt.

(ii) Eine Teilmenge U eines metrischen Raumes heisst *offen*, falls $\forall x \in U$ ein $\epsilon > 0$ existiert so, dass die “ ϵ -Kugel um x ” voll in U enthalten ist, d.h.

$$\exists \epsilon > 0 : B_\epsilon(x) = \{y \in X \mid d(x, y) < \epsilon\} \subset U \quad (15.35)$$

· Insbesondere sind für $x \in X$ und $R > 0$ die offenen Kugeln (vgl. (7.17)):

$$B_R(x) = \{y \in X \mid d(x, y) < R\} \quad (15.36)$$

offen in X : Für $y \in B_R(x)$ gilt per Definition $R - d(x, y) > 0$ und es existiert noch ein $\epsilon > 0$ mit $\epsilon < R - d(x, y)$. Aus der Dreiecksungleichung folgt dann $\forall z \in B_\epsilon(y)$: $d(z, x) \leq d(z, y) + d(y, x) < \epsilon + d(x, y) < R$, d.h. also $B_\epsilon(y) \subset B_R(x)$.

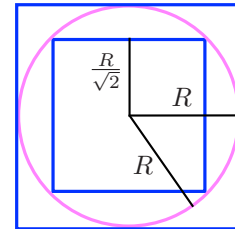
· Typische Beispiele abgeschlossener Mengen sind die “abgeschlossenen Kugeln”

$$\overline{B_R(x)} = \{y \in X \mid d(x, y) \leq R\} \quad (15.37)$$

Bew.: Es sei $(y_k) \subset \overline{B_R(x)}$ konvergent gegen $y \in X$. Dann gilt $\forall \epsilon > 0$: $d(x, y) \leq d(x, y_k) + d(y_k, y) \leq R + \epsilon$ für k gross genug. Es folgt also $d(y, x) \leq R + \epsilon \forall \epsilon > 0$, und dies geht nur wenn $d(y, x) \leq R$.

· Die offenen/abgeschlossenen Kugeln von $\mathbb{R}$ mit der üblichen Abstandsfunktion sind genau die (beschränkten) offenen/abgeschlossenen Intervalle. Es gibt aber kein gutes allgemeines Analogon der halboffenen Intervalle.

· Für $X = \mathbb{R}^n$ und $d = d_2$ aus (15.6) sind die $\overline{B_R}$ und B_R einfach die euklidischen Vollkugeln mit bzw. ohne die Kugelfläche. Für $d = d_\infty$ aus (15.8) sind es Quader der Kantenlänge $2R$. Die Äquivalenz der Normen 15.16 besagt, dass in jede Kugel noch ein Quader passt und umgekehrt.



· In einem metrischen Raum vom Typ (15.24) ist $\forall x \in X$:

$$B_R(x) = \begin{cases} X & \text{falls } R > 1 \\ \{x\} & \text{falls } R \leq 1 \end{cases} \quad \overline{B_R(x)} = \begin{cases} X & \text{falls } R \geq 1 \\ \{x\} & \text{falls } R < 1 \end{cases} \quad (15.38)$$

⁸⁶man denke sich dazu: unter Folgenkonvergenz

§ 15. METRISCHE RÄUME

· Man beachte, dass “offen” und “abgeschlossen” sich nicht gegenseitig ausschliessen. Insbesondere sind X und die leere Menge (mit den üblichen Verabredungen) beides, so dass gilt:

Lemma 15.19. $U \subset X$ ist offen $\Leftrightarrow X \setminus U$ ist abgeschlossen.

Beweis. Für $U = \emptyset$ oder $X \setminus U = \emptyset$ ist die Aussage wahr. Andernfalls:

“ $\Rightarrow$ ”: Sei $(x_k) \subset X \setminus U$ eine in X konvergente Folge mit Grenzwert $x = \lim x_k \in X$. Wäre $x \in U$ so gäbe es ein $\epsilon_* > 0$ so dass $B_{\epsilon_*}(x) \subset U$. Wegen der Konvergenz der Folge gibt es ein K_{ϵ_*} so dass $d(x, x_k) < \epsilon_*$, d.h. $x_k \in B_{\epsilon_*}(x) \forall k \geq K_{\epsilon_*}$. Dies steht im Widerspruch zu $x_k \in X \setminus U$. Es muss also $x \in X \setminus U$ gelten, und daher ist $X \setminus U$ abgeschlossen.

“ $\Leftarrow$ ”: Sei $x \in U$. Gäbe es kein $\epsilon > 0$ mit $B_\epsilon(x) \subset U$, d.h. wäre für jedes $\epsilon > 0$ $B_\epsilon(x) \cap (X \setminus U) \neq \emptyset$, so gäbe es insbesondere für jedes $k \in \mathbb{N}$ ein $x_k \in B_{1/k}(x) \cap (X \setminus U)$. Die Folge $(x_k) \subset X \setminus U$ konvergierte dann gegen $x \in U$, im Widerspruch zur Abgeschlossenheit von $X \setminus U$. Es muss also ein $\epsilon > 0$ geben, für das $B_\epsilon(x) \subset U$, und daher ist U offen. $\square$

Man kann auch die Offenheit einer Teilmenge durch die Konvergenzeigenschaften von Folgen charakterisieren:

Lemma 15.20. $U \subset X$ ist offen $\Leftrightarrow$ Für jede Folge $(x_k) \subset X$, welche gegen ein $x \in U$ konvergiert, liegen fast alle x_k in U . (Fast alle heisst: “bis auf endlich viele”.)

Beweis. “ $\Rightarrow$ ”: Sei $\epsilon > 0$ so, dass $B_\epsilon(x) \subset U$, und K_ϵ so, dass $d(x, x_k) < \epsilon$ falls $k \geq K_\epsilon$. Dann ist $x_k \in B_\epsilon(x) \subset U \forall k \geq K_\epsilon$. “ $\Leftarrow$ ”: Wäre U nicht offen, so gäbe es für ein $x \in U$ für jedes $k \in \mathbb{N}$ ein $x_k \in B_{1/k}(x) \setminus U$. Es gilt $x_k \rightarrow x$ aber kein x_k liegt in U . $\square$

Ist der umgebende metrische Raum fixiert, so sagt man statt “offene/abgeschlossene Teilmenge” auch “offene/abgeschlossene Menge”.

Proposition 15.21. (i) Ist I eine beliebige Menge, und $(U_i)_{i \in I}$ eine durch I indizierte Familie offener Mengen, so ist ihre Vereinigung

$$\bigcup_{i \in I} U_i \quad (15.39)$$

wieder offen.

(ii) Ist E eine endliche Menge, und $(U_i)_{i \in E}$ eine durch E indizierte Familie offener Mengen, so ist ihr Durchschnitt

$$\bigcap_{i \in E} U_i \quad (15.40)$$

wieder offen.

Komplementär dazu ist die endliche Vereinigung und der beliebige Durchschnitt abgeschlossener Mengen wieder abgeschlossen.

Beweis. (i) Für beliebiges $x \in V := \bigcup_{i \in I} U_i$ existiert mindestens ein $i_x \in I$ so dass $x \in U_{i_x}$. Da U_{i_x} offen ist, existiert ein $\epsilon_{i_x} > 0$ so dass $B_{\epsilon_{i_x}}(x) \subset U_{i_x}$. Es folgt $B_{\epsilon_{i_x}}(x) \subset V$.

(ii) Für jedes $x \in D := \bigcap_{i \in E} U_i$ gilt: $x \in U_i \forall i \in E$. Da die U_i offen sind, existiert daher für jedes $i \in E$ ein $\epsilon_i > 0$ so dass $B_{\epsilon_i}(x) \subset U_i$. Mit $\epsilon = \min\{\epsilon_i \mid i \in E\} > 0$ gilt dann $B_\epsilon(x) \subset U_i \forall i \in E$, also $B_\epsilon(x) \subset D$.

Die Aussagen für abgeschlossene Mengen überlegt man sich entweder direkt oder durch Anwendung der DeMorganschen Gesetze ($X \setminus (A \cup B) = (X \setminus A) \cap (X \setminus B)$, $X \setminus (A \cap B) = (X \setminus A) \cup (X \setminus B)$ etc.) und 15.19. $\square$

Der beliebige Durchschnitt offener Mengen ist nicht notwendig offen. Erkläre wo der Beweis schief geht, und gib ein Gegenbeispiel!

Lemma 15.22. *Eine abgeschlossene Teilmenge eines vollständigen metrischen Raumes ist vollständig (als metrischer Raum für sich).*

Beweis. Sei (X, d_X) ein vollständiger metrischer Raum, und $A \subset X$ eine abgeschlossene Teilmenge. Die Abstandsfunktion auf A ist $d_A = d_X|_{A \times A}$ (s. (15.34)). Ist daher $(x_k) \subset A$ eine Cauchy-Folge bzgl. d_A , so ist sie trivialerweise auch als Folge in X bezüglich d_X eine Cauchy-Folge. Wegen der Vollständigkeit von (X, d_X) ist (x_k) konvergent in X und wegen der Abgeschlossenheit von A liegt ihr Grenzwert in A . Jede Cauchy-Folge in (A, d_A) konvergiert also in A , d.h. (A, d_A) ist vollständig. $\square$

Definition 15.23. Wir nennen einen metrischen Raum (X, d) beschränkt, falls die Teilmenge

$$\{d(x, y) \mid x, y \in X\} \quad (15.41)$$

von $\mathbb{R}$ beschränkt ist, m.a.W. falls eine Zahl $M > 0$ existiert so dass $d(x, y) \leq M \forall x, y \in X$. Falls $X \neq \emptyset$, so heisst das Supremum der Menge (15.41) der Durchmesser von X , geschrieben $\text{diam } X$.

· Die Begriffe “offene und abgeschlossene Mengen” und “Durchmesser” hängen nicht von der Vollständigkeit des betreffenden metrischen Raums ab.

· Wir wenden den Begriff des Durchmessers auch auf (nicht-leere) Teilmengen von X mit der eingeschränkten Abstandsfunktion an. (Man beachte dabei, dass der Durchmesser nicht angenommen werden muss, z.B. für offene Kugeln im $\mathbb{R}^n$.) Damit erhalten wir eine nützliche Verallgemeinerung des ”Intervallschachtelungsprinzips” Thm. 14.16:

Proposition 15.24. *Sei $(A_k)_{k \in \mathbb{N}} \subset \mathcal{P}(X)$ eine Folge von Teilmengen eines vollständigen metrischen Raumes (X, d) derart, dass gilt:*

(i) $\forall k \in \mathbb{N}$ ist A_k nicht-leer, abgeschlossen und beschränkt.

(ii) $A_{k+1} \subset A_k \forall k$

(iii) $\text{diam } A_k \rightarrow 0$ für $k \rightarrow \infty$.

Dann existiert genau ein $x \in \bigcap_{l=1}^{\infty} A_l$.

Beweis. Wähle für jedes $k \in \mathbb{N}$ ein $x_k \in A_k$. Wegen der Schachtelung ist (x_k) eine Cauchy-Folge, und für jedes $l \in \mathbb{N}$ liegen fast alle Glieder von (x_k) in A_l . Wegen der Abgeschlossenheit von A_l liegt dann der Grenzwert $x = \lim x_k$ in jedem A_l , also in $\bigcap_{l=1}^{\infty} A_l$. Ist y ein (a priori) anderer Punkt in $\bigcap_{l=1}^{\infty} A_l$, so folgt $d(x, y) < \epsilon \forall \epsilon > 0$, d.h. $y = x$. $\square$

§ 15. METRISCHE RÄUME

Zum Schluss kehren wir zu $\mathbb{R}^n$ zurück und beweisen die folgende Verallgemeinerung des Satzes von Bolzano-Weierstrass (Thm. 14.17)

Proposition 15.25. *Jede beschränkte Folge $(x_k) \subset \mathbb{R}^n$ besitzt eine konvergente Teilfolge.*

Bemerkung. Wir machen diese Behauptung für jede Norm auf $\mathbb{R}^n$. Zum Beweis zeigen wir die Aussage zunächst für die Norm $\|\cdot\|_\infty$ unter Benutzung der bereits bewiesenen Vollständigkeit von $(\mathbb{R}^n, \|\cdot\|_\infty)$ (vgl. Prop. 15.12). Anschliessend benutzen wir dieses Resultat, um zu zeigen, dass alle Normen auf $\mathbb{R}^n$ äquivalent sind, d.h. Thm. 15.17. Daraus folgt dann mit Argumenten wie bei Lemma 15.15, dass 15.25 mit jeder Norm gilt.

Beweis von 15.25 für $\|\cdot\|_\infty$. In Imitation des Beweises von 14.17 konstruieren wir zunächst eine Folge $(A_l)_{l \in \mathbb{N}}$ wie in 15.24 mit der Eigenschaft, dass jedes A_l unendlich viele x_k enthält: Da $\{x_k \mid k \in \mathbb{N}\}$ beschränkt ist, existiert ein $R > 0$ so, dass $\|x_k\|_\infty \leq R \forall k$, d.h.,

$$(x_k) \subset \overline{B_R(0)} = \{y \in \mathbb{R}^n \mid \|y\|_\infty \leq R\} = [-R, R]^n \quad (15.42)$$

Laut (15.37) sind die $\overline{B_R(0)}$ (geometrisch gesehen Quader) abgeschlossen. Im ersten Schritt teilen wir $\overline{B_R(0)}$ in 2^n Teile,

$$\overline{B_{R/2}((\pm \frac{R}{2}, \pm \frac{R}{2}, \dots, \pm \frac{R}{2})^T)} \quad (15.43)$$

vom Durchmesser R , und wählen für A_1 einen dieser Quader, welcher unendlich viele Folgenglieder enthält. Wir verfahren analog im l -ten Schritt, erhalten eine Schachtelung mit den gewünschten Eigenschaften, und setzen $\{x\} = \cap_{l=1}^\infty A_l$. Eine Wahl $x_{k_l} \in A_l$ so, dass (k_l) streng monoton wächst, liefert eine gegen x konvergente Teilfolge (x_{k_l}) . $\square$

Bemerkung. Für eine unendliche Menge, ausgerüstet mit der Abstandsfunktion (15.24) ist die Aussage offenbar falsch: Wegen der Unendlichkeit von X existiert eine injektive Folge, welche in der diskreten Metrik keinen Häufungswert hat, obwohl X beschränkt (und auch vollständig) ist. Der Beweis scheitert daran, dass man für $R < 1$ unendlich viele Kugeln vom Radius R braucht, um X zu überdecken, vgl. (15.38).

Proposition 15.26. *Sei $\|\cdot\|$ eine beliebige Norm auf $\mathbb{R}^n$. Dann existieren Konstanten $c > 0$, $C > 0$ so, dass für alle $x \in \mathbb{R}^n$*

$$c \cdot \|x\|_\infty \leq \|x\| \leq C \cdot \|x\|_\infty \quad (15.44)$$

Beweis. Sind für $i = 1, \dots, n$ $e_i = (0, \dots, \overset{i\text{-te Stelle}}{1}, \dots, 0)^T$ die Elemente der Standardbasis des $\mathbb{R}^n$, so folgt für $x = \sum_{i=1}^n x^i e_i$ mit Hilfe der Dreiecksungleichung und der Homogenität von $\|\cdot\|$:

$$\|x\| \leq \sum_{i=1}^n |x^i| \cdot \|e_i\| \leq n \cdot \max\{\|e_i\|\} \cdot \|x\|_\infty \quad (15.45)$$

Damit haben wir C gefunden, und nach den Bemerkungen zu Lemma 15.15 folgt bereits, dass Konvergenz in $\|\cdot\|_\infty$ Konvergenz in $\|\cdot\|$ impliziert. (Aus (15.45) folgt nämlich, dass die ϵ -Kugel bzgl. $\|\cdot\|$ die ϵ/C -Kugel bzgl. $\|\cdot\|_\infty$ enthält. Eine Folge läuft daher in der ϵ -Kugel bzgl. $\|\cdot\|$ sobald sie in die ϵ/C -Kugel bzgl. $\|\cdot\|_\infty$ hineingelaufen ist.)

Den Beweis der Existenz von c führen wir indirekt. Wir nehmen also an, es gäbe kein c wie verlangt. (Dies bedeutet anschaulich, dass eine gewisse (oder wegen der Homogenität der Normen dazu äquivalent, jede) Kugel bzgl. $\|\cdot\|_\infty$ keine (noch so kleine) Kugel bzgl. $\|\cdot\|$ enthält.) Dann gibt es insbesondere für jedes k ein $\tilde{x}_k \in \mathbb{R}^n$ mit $\|\tilde{x}_k\| < \frac{1}{k} \|\tilde{x}_k\|_\infty$. Es folgt $\tilde{x}_k \neq 0$ und mit $x_k := \tilde{x}_k / \|\tilde{x}_k\|_\infty$ erhalten wir eine Folge (x_k) mit

$$\|x_k\| = \frac{\|\tilde{x}_k\|}{\|\tilde{x}_k\|_\infty} < \frac{1}{k} \quad \text{und} \quad \|x_k\|_\infty = 1 \quad (15.46)$$

(Die Existenz von (x_k) besagt, dass keine der $1/k$ -Kugeln bzgl. $\|\cdot\|$ in der (abgeschlossenen) Einheitskugel bzgl. $\|\cdot\|_\infty$ liegt.) Die Folge (x_k) ist bzgl. $\|\cdot\|_\infty$ beschränkt, und besitzt daher nach 15.25 eine konvergente Teilfolge (x_{k_l}) . Deren Grenzwert x erfüllt noch $\|x\|_\infty = 1$. Andererseits folgt aus obiger Konvergenzimplikation, dass auch $\|x - x_{k_l}\| \rightarrow 0$ für $l \rightarrow \infty$, und dann aus

$$\|x\| \leq \|x - x_{k_l}\| + \|x_{k_l}\| \rightarrow 0 \quad \text{für } l \rightarrow \infty \quad (15.47)$$

dass $\|x\| = 0$, d.h. $x = 0$, ein Widerspruch zur Positivität der Norm $\|\cdot\|$. $\square$

Da jeder endlich-dimensionale Vektorraum durch Auszeichnung einer Basis isomorph zum $\mathbb{R}^n$ mit der Standardbasis ist,⁸⁷ haben wir zusammengefasst sowohl Thm. 15.17 als auch Prop. 15.25 vollständig bewiesen. $\square$

Auf unendlich-dimensionalen reellen Vektorräumen gibt es inäquivalente Normen und verschiedene Vollständigkeitsbegriffe. Insbesondere existieren auch unvollständige (unendlich-dimensionale) normierte Vektorräume, und beschränkte Folgen ohne konvergente Teilfolgen.

Als möglicherweise aufschlussreiches endlich-dimensionales Beispiel betrachten wir auf $V = \mathbb{Q}^2 = \{(x, y) \mid x, y \in \mathbb{Q}\}$ die Abbildung

$$\begin{aligned} \|\cdot\|_\star : \mathbb{Q}^2 &\rightarrow \mathbb{R} \\ (x, y) &\mapsto \|(x, y)\|_\star := |x - \sqrt{2}y| \end{aligned} \quad (15.48)$$

wo $|\cdot|$ den gewöhnlichen Absolutbetrag auf $\mathbb{R}$ bezeichnet. Diese Funktion ist offensichtlich (absolut) homogen und erfüllt die Dreiecksungleichung. Positivität folgt aus der Tatsache, dass die Gleichung $x = \sqrt{2}y$ über $\mathbb{Q}$ keine nicht-triviale Lösung hat.

Ob man $\|\cdot\|_\star$ als Norm gelten lässt, hängt etwas davon ab, wie man den Begriff der Def. 15.2 auf Vektorräume über beliebigen angeordneten Körpern verallgemeinert. Im Sinne der Physik würde ich darauf bestehen wollen, dass in solchen Fällen

⁸⁷Übungsaufgabe: Man zeige, dass ein solcher Isomorphismus mit der Äquivalenz von Normen verträglich ist.

§ 16. REIHEN

Normen stets (nicht-negative) Werte im Grundkörper annehmen, unabhängig davon, ob dieser grösser oder kleiner als $\mathbb{R}$ ist.⁸⁸ Für die Algebra ist dies nicht unbedingt zweckmässig. Hier werden schon auf den Körpern selber (nur) Absolutbeträge mit Werten in $\mathbb{R}$ zugelassen, damit sie sich in sinnvoller Weise als metrische Räume vervollständigen lassen, und dies überträgt sich dann auf Vektorräume.

In jedem Fall aber induziert $\|\cdot\|_\star$ eine Abstandsfunktion und macht damit $\mathbb{Q}^2$ zu einem metrischen Raum (über $\mathbb{R}$) im Sinne von Def. 15.1. Dieser Raum ist unvollständig. Interessanter jedoch ist, dass zwar

$$\|(x, y)\|_\star \leq 2(|x| + |y|) = 2\|(x, y)\|_1, \quad (15.49)$$

andererseits aber für eine Folge (x_k) rationaler Zahlen, die (in $\mathbb{R}$) gegen $\sqrt{2}$ konvergiert,

$$\begin{aligned} \lim_{k \rightarrow \infty} \|(x_k, 1)\|_\star &= 0 \\ \text{aber } \lim_{k \rightarrow \infty} \|(x_k, 1)\|_1 &= \sqrt{2} + 1 \end{aligned} \quad (15.50)$$

Daraus folgt, dass kein $c > 0$ existiert so dass $c\|(x, y)\|_1 \leq \|(x, y)\|_\star \quad \forall (x, y) \in \mathbb{Q}^2$, d.h. $\|\cdot\|_\star$ ist echt schwächer als $\|\cdot\|_1$!

§ 16 Reihen⁸⁹

Reihen, intuitiv aufgefasst als unendliche Summen der Form

$$\sum_{k=0}^{\infty} a_k = a_0 + a_1 + a_2 + \cdots \quad (16.1)$$

sind ein wichtiger Fall der Vorstellung mathematischer Objekte durch Grenzprozesse, sowohl aus praktischer als auch aus historischer Sicht. Formal sind sie nichts anderes als Folgen (s_n) in normierten Vektorräumen, bei denen man das Augenmerk auf die Zuwächse

$$a_k = s_k - s_{k-1} \quad (16.2)$$

anstatt auf die s_n selber richtet, m.a.W.:

Definition 16.1. Sei $(V, \|\cdot\|)$ ein normierter Vektorraum (standardmässig: über $\mathbb{R}$, ggfs. $\mathbb{C}$). Für eine beliebige Folge $(a_k)_{k \in \mathbb{N}_0} \subset V$ heisst die Folge $(s_n)_{n \in \mathbb{N}_0}$ mit

$$\begin{aligned} s_0 &:= a_0 \\ s_1 &:= s_0 + a_1 = a_0 + a_1 \\ s_2 &:= s_1 + a_2 = a_0 + a_1 + a_2 \\ &\vdots \\ s_n &:= s_{n-1} + a_n = \sum_{k=0}^n a_k \\ &\vdots \end{aligned} \quad (16.3)$$

⁸⁸Das Beispiel (15.48) erfüllt diese Bedingung offenbar nicht.

⁸⁹Dieser § ist von der Themenwahl und -anordnung stark an Kapitel 6 der Analysis I von Königsberger angelehnt. Mit der Darstellung bin ich aber aktuell nicht mehr sehr glücklich.

die der Folge (a_k) zugeordnete *unendliche Reihe*, oder kurz Reihe. Die a_k heissen *Glieder* der Reihe, die s_n die *Partialsummen*.

• Die Reihe heisst konvergent, falls die Folge der Partialsummen konvergiert (im Sinne der Defs. 15.8, 15.9). Man schreibt

$$\sum_{k=0}^{\infty} a_k \quad \text{oder auch} \quad \sum a_k \quad (16.4)$$

sowohl für die Folge der Partialsummen als auch für ihren Grenzwert, falls dieser existiert. Dieser Grenzwert heisst dann auch *Wert der Reihe* oder auch *Summe der Folge* (a_k) .

In dieser Definition steht $\mathbb{N}_0 := \mathbb{N} \cup \{0\}$ für die Menge der nicht-negativen ganzen (a.k.a., “noch-natürlichen”) Zahlen mit ihrer gewöhnlichen Anordnung. Dies ist der häufigste Definitionsbereich des Laufindex, k . Wie bei Folgen ändern sich die Konvergenz einer Reihe nicht, wenn man *endlich viele* Glieder weglässt, hinzufügt (z.B., mit negativem Index), oder abändert. Der Wert einer konvergenten Reihe bleibt dabei im Allgemeinen natürlich nicht gleich.

Lemma 16.2. *Die Glieder einer konvergenten Reihe bilden eine Nullfolge (im Sinne der (Verallgemeinerung von) Def. 14.6).*

Beweis. Folgt unmittelbar aus $a_k = s_k - s_{k-1}$ und den (auf normierte Vektorräume ausgedehnte) Rechenregeln für Folgen, speziell 14.10 (α) und Lemma 14.11:

$$\lim_{k \rightarrow \infty} a_k = \lim_{n \rightarrow \infty} s_n - \lim_{n \rightarrow \infty} s_{n-1} = \sum_{k=0}^{\infty} a_k - \sum_{k=0}^{\infty} a_k = 0 \quad (16.5)$$

□

(Die Umkehrung gilt natürlich wohlgemerkt nicht: Eine Reihe kann divergieren, auch wenn die a_k eine Nullfolge bilden. Beispiele folgen.) Unser Ziel in diesem § ist eine Übersicht über Konvergenzkriterien für Reihen sowie die Feststellung einiger Rechenregeln. Wir behandeln zunächst komplexe Reihen, mit einigen anordnungs-spezifischen Bemerkungen im rein reellen Fall. Abschliessend und als Vorbereitung für alles Weitere richten wir den Blick auf die für die Physik wichtigen Potenzreihen.

Lemma 16.3. *Sind $\sum a_k$, $\sum b_k$ zwei konvergente Reihen, und $\alpha \in \mathbb{R}$ (ggfs., $\mathbb{C}$), so ist auch $\sum (a_k + \alpha b_k)$ konvergent und es gilt $\sum (a_k + \alpha b_k) = \sum a_k + \alpha \sum b_k$.*

Beweis. Für die Folgen der Partialsummen gilt wegen der Assoziativität und Kommutativität bei endlichen Summen

$$\sum_{k=0}^n (a_k + \alpha b_k) = \sum_{k=0}^n a_k + \alpha \sum_{k=0}^n b_k \quad (16.6)$$

und daher folgt die Behauptung unmittelbar aus den Rechenregeln für Folgen. □

§ 16. REIHEN

Konvergente Reihen bilden also, ebenso wie konvergente Folgen, mit den offensichtlichen Rechenoperationen einen (unendlich-dimensionalen) reellen (oder komplexen) Vektorraum. Die interessanten Fragen beim Rechnen mit Reihen sind die Abhängigkeit der Konvergenz von der Variation unendlich vieler Reihenglieder sowie die Unabhängigkeit vom Vertauschen von Reihengliedern, s. Def. 16.11.

Beispiel 16.4. Für $z \in \mathbb{C}$ ist die *geometrische Reihe*

$$\sum_{k=0}^{\infty} z^k = \left(s_n = \sum_{k=0}^n z^k \right)_{n=0,1,\dots} \stackrel{\text{falls } z \neq 1}{=} \left(\frac{1 - z^{n+1}}{1 - z} \right)_{n=0,1,\dots} \quad (16.7)$$

- für $|z| < 1$ konvergent mit Wert $\sum_{k=0}^{\infty} z^k = \frac{1}{1-z}$. (Dies folgt aus der expliziten Form für die Partialsummen und $z^n \rightarrow 0$ für $n \rightarrow \infty$.)
- sonst divergent ((z^k) ist keine Nullfolge für $|z| \geq 1$).

Beispiel 16.5. Für $s \in \mathbb{Q}$ ist die Dirichlet-Riemann-Reihe

$$\sum_{k=1}^{\infty} \frac{1}{k^s} = \left(\text{Es gibt im Allgemeinen keine geschlossene Form für die Partialsummen.} \right) \quad (16.8)$$

- für $s > 1$ konvergent.
- für $s \leq 1$ divergent.

Bew.: Da alle Glieder positiv sind, wächst die Folge der Partialsummen streng monoton.

- Für $s > 1$, $n \in \mathbb{N}$ sei l so, dass $2^l > n$. Dann gilt

$$\begin{aligned} \sum_{k=1}^n \frac{1}{k^s} &\leq \sum_{k=1}^{2^l-1} \frac{1}{k^s} \\ &= 1 + \underbrace{\frac{1}{2^s} + \frac{1}{3^s}}_{\leq 2 \cdot \frac{1}{2^s}} + \underbrace{\frac{1}{4^s} + \dots + \frac{1}{7^s}}_{\leq 2^2 \cdot \frac{1}{2^{2s}}} + \dots + \underbrace{\frac{1}{2^{(l-1)s}} + \dots + \frac{1}{(2^l-1)^s}}_{\leq 2^{l-1} \cdot \frac{1}{2^{(l-1)s}}} \quad (16.9) \\ &\leq 1 + \frac{1}{2^{s-1}} + \left(\frac{1}{2^{s-1}} \right)^2 + \dots + \left(\frac{1}{2^{s-1}} \right)^{l-1} \\ &\leq \frac{1}{1 - \frac{1}{2^{s-1}}} \quad (\text{geometrische Reihe; } 2^{1-s} < 1 \text{ für } s > 1) \end{aligned}$$

Die Folge der Partialsummen ist also monoton wachsend und nach oben beschränkt, also konvergent nach Theorem 14.13.

· Für $s \leq 1$ gilt für alle k : $\frac{1}{k^s} \geq \frac{1}{k}$.⁹⁰ Dann gilt für alle $l \in \mathbb{N}$ und $n \geq 2^l$

$$\begin{aligned} \sum_{k=1}^n \frac{1}{k^s} &\geq \sum_{k=1}^{2^l} \frac{1}{k^s} \\ &= 1 + \frac{1}{2} + \underbrace{\frac{1}{3} + \frac{1}{4}}_{\geq 2 \cdot \frac{1}{4}} + \underbrace{\frac{1}{5} + \cdots + \frac{1}{8}}_{\geq 4 \cdot \frac{1}{8}} + \cdots + \underbrace{\frac{1}{2^{l-1}+1} + \cdots + \frac{1}{2^l}}_{\geq 2^{l-1} \cdot \frac{1}{2^l}} \quad (16.10) \\ &= 1 + \frac{l}{2} \rightarrow \infty \quad \text{für } l \rightarrow \infty \end{aligned}$$

Insbesondere ist die harmonische Reihe $\sum k^{-1}$ divergent.

Diese Beispiele illustrieren bereits die zwei wesentlichen Ideen zum Nachweis von Konvergenz/Divergenz von Reihen: (i) Vergleich mit bekannten Reihen, und (ii) Monotonie für Reihen mit positiven Gliedern. Eine grundsätzliche Überlegung ist

Lemma 16.6. Ist $(a_k) \subset \mathbb{C}$ und $(c_k) \subset \mathbb{R}_{\geq 0}$ mit:

1. $|a_k| \leq c_k \quad \forall k$
2. $\sum c_k$ ist konvergent.

Dann gilt: $\sum a_k$ ist konvergent und $\left| \sum_{k=0}^{\infty} a_k \right| \leq \sum_{k=0}^{\infty} c_k$. Man nennt eine solche Reihe $\sum c_k$ eine konvergente Majorante von $\sum a_k$.

Beweis. Man prüft das Cauchy-Kriterium für $s_n := \sum_{k=0}^n a_k$, gestützt auf dasjenige

für $u_n := \sum_{k=0}^n c_k$: Für $\epsilon > 0$ sei N_ϵ so, dass $|u_n - u_m| < \epsilon$ für $n, m \geq N_\epsilon$. Dann gilt für solche n, m , oBdA mit $n \geq m$:

$$|s_n - s_m| = \left| \sum_{k=m+1}^n a_k \right| \leq \sum_{k=m+1}^n |a_k| \leq \sum_{k=m+1}^n c_k = |u_n - u_m| < \epsilon \quad (16.11)$$

Daraus folgt mit Theorem 14.21 die Konvergenz der Reihe. Die Abschätzung für ihren Wert folgt aus $\left| \sum_{k=0}^n a_k \right| \leq \sum_{k=0}^n c_k$ zusammen mit den Rechenregeln für Folgen. $\square$

Definition 16.7. Eine Reihe $\sum a_k$ heisst absolut konvergent, falls die Reihe der Absolutbeträge $\sum |a_k|$ konvergiert.

Proposition 16.8. Eine absolut konvergente Reihe ist konvergent.

Beweis. Die Reihe der Absolutbeträge ist eine konvergente Majorante, also folgt die Aussage direkt aus Lemma 16.6. $\square$

⁹⁰Dies folgt aus der (trivialen) Abschätzung $x^r \geq 1$ falls $x > 1$ und $r > 1$ via $x = k$, $r = 1 - s$.

§ 16. REIHEN

Die Umkehrung von 16.8 gilt aber nicht, wie das folgende Beispiel zeigt:

Beispiel 16.9. Die alternierende harmonische Reihe,

$$\sum_{k=1}^{\infty} \frac{(-1)^{k-1}}{k} = 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \cdots \quad (16.12)$$

konvergiert, aber nicht absolut.

Hierbei folgt die Divergenz der Reihe der Absolutbeträge aus (16.10), die Konvergenz von (16.12) aus dem sog. Leibniz-Kriterium für alternierende Reihen:

Lemma 16.10. *Sei (a_k) eine monoton fallende (oder auch wachsende) Nullfolge reeller Zahlen. Dann konvergiert*

$$\sum_{k=1}^{\infty} (-1)^k a_k \quad (16.13)$$

Beweis. Sei $s_n = \sum_{k=1}^n (-1)^k a_k$ die Folge der Partialsummen. Für gerade $n = 2m$ gilt $s_{2m} = s_{2m-1} + a_{2m} \geq s_{2m-1}$ sowie $s_{2m+1} = s_{2m-1} + a_{2m} - a_{2m+1} \geq s_{2m-1}$ (Monotonie!) und $s_{2m+2} = s_{2m} - a_{2m+1} + a_{2m+2} \leq s_{2m}$, d.h.

$$[s_{2m+1}, s_{2m+2}] \subset [s_{2m-1}, s_{2m}] \quad \text{und} \quad s_{2m} - s_{2m-1} = a_{2m} \rightarrow 0 \quad (16.14)$$

Die $([s_{2m-1}, s_{2m}])_{m=1,2,\dots}$ bilden also eine Intervallschachtelung gemäss Def. 14.15, und ihr Durchschnitt ist gleich dem Grenzwert der Randpunkte gemäss Theorem 14.16,

$$\bigcap_{m=1}^{\infty} [s_{2m-1}, s_{2m}] = \left\{ \lim s_{2m} = \lim s_{2m-1} = \sum_{k=1}^{\infty} (-1)^k a_k \right\} \quad (16.15)$$

□

Absolute Konvergenz

Witzigerweise klingt im Adjektiv “absolut” noch die Bedeutung an, dass man die Glieder einer absolut konvergenten Reihe beliebig umordnen oder umgruppieren kann, ohne die Konvergenz oder den Wert der Reihe zu beeinflussen, während dies für konvergente, aber nicht absolut konvergente Reihen nicht unbedingt der Fall ist.

Definition 16.11. Unter einer *Umordnung* einer Reihe $\sum_{k=1}^{\infty} a_k$ verstehen wir eine Reihe $\sum_{l=1}^{\infty} b_l$, die aus $\sum a_k$ durch Angabe einer *bijektiven Abbildung* $\mathbb{N} \rightarrow \mathbb{N}$, $l \mapsto k_l$ via $b_l := a_{k_l}$ hervorgeht. (Man plant also das Absummieren *aller* Glieder von $\sum a_k$, nur eben in einer anderen *Reihenfolge*, vgl. aber auch mit dem Begriff einer Teilfolge in Def. 14.3.)

Ordnen wir beispielsweise die alternierende harmonische Reihe (16.12) so um, dass auf ein positives stets zwei negative Glieder folgen, d.h.

$$1 - \frac{1}{2} - \frac{1}{4} + \frac{1}{3} - \frac{1}{6} - \frac{1}{8} + \cdots + \frac{1}{2m-1} - \frac{1}{2(2m-1)} - \frac{1}{4m} + \cdots \quad (16.16)$$

In Worten: Die m -te Gruppe enthält mit entsprechenden Vorzeichen die Kehrwerte der m -ten ungeraden Zahl, sowie der $(2m - 1)$ -ten und $2m$ -ten geraden. In der umgeordneten Reihe sind dies die Glieder Nummer $l = 3m - 2$, $3m - 1$ und $3m$. Die umordnende Bijektion ist also

$$k_l = \begin{cases} \frac{2l+1}{3} & l \equiv 1 \pmod{3} \\ 2 \frac{2l-1}{3} & l \equiv 2 \pmod{3} \\ 4 \frac{l}{3} & l \equiv 0 \pmod{3} \end{cases} \quad (16.17)$$

Dann gilt wegen $\frac{1}{2m-1} - \frac{1}{2(2m-1)} - \frac{1}{4m} = \frac{1}{2} \left(\frac{1}{2m-1} - \frac{1}{2m} \right)$, dass die Summe der ersten $3n$ Glieder von (16.16) gleich

$$\frac{1}{2} \left(1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \cdots + \frac{1}{2n-1} - \frac{1}{2n} \right) \quad (16.18)$$

d.h. der Hälfte der Summe der ersten $2n$ Glieder der alternierenden harmonischen Reihe ist. Die dazwischenliegenden Partialsummen liegen nicht weiter als $\frac{1}{4n}$ bzw. $\frac{1}{2n+1}$ davon entfernt, so dass (16.16) gegen die Hälfte von (16.12) konvergiert.

Unter absoluter Konvergenz (äquivalent dazu: bei Vorhandensein einer konvergenten Majorante) kann so etwas nicht passieren.

Proposition 16.12. *Sei $\sum_{k=1}^{\infty} a_k$ eine absolut konvergente Reihe. Dann konvergiert jede Umordnung von $\sum a_k$ ebenfalls absolut und hat den gleichen Wert.*

Wir benutzen als “Trivialkriterium” hierfür:

Lemma 16.13. *Eine Reihe $\sum a_k$ ist genau dann absolut konvergent, wenn die Menge*

$$\mathcal{F} = \left\{ \sum_{k \in J} |a_k| \mid J \subset \mathbb{N} \text{ endliche Teilmenge} \right\} \quad (16.19)$$

nach oben beschränkt ist. In diesem Fall gilt $\sum_{k=1}^{\infty} |a_k| = \sup \mathcal{F}$.

Beweis. Ist die Reihe absolut konvergent, so ist der Wert von $\sum |a_k|$ klarerweise eine obere Schranke für $\mathcal{F}$: Für jede endliche Teilmenge $J \subset \mathbb{N}$ gibt es ein N_J so dass $J \subset \{1, \dots, N_J\}$. Dann ist

$$\begin{array}{c} \text{Monotonie von } \sum |a_k| \\ \downarrow \\ \sum_{k=1}^{\infty} |a_k| \geq \sum_{k=1}^{N_J} |a_k| \geq \sum_{k \in J} |a_k| \end{array} \quad (16.20)$$

Ist umgekehrt $\mathcal{F}$ nach oben beschränkt, so wächst $(\sum_{k=1}^n |a_k|)_{n=1,2,\dots}$ monoton und ist nach oben beschränkt, also konvergent nach Theorem 14.13.

· Um noch zu zeigen, dass in diesem Fall $\sum |a_k| = \sup \mathcal{F}$, bemerken wir, dass für jedes $\epsilon > 0$ wegen der Konvergenz von $\sum |a_k|$ ein N_ϵ existiert so dass $\sum_{k=1}^{N_\epsilon} |a_k| > \sum_{k=1}^{\infty} |a_k| - \epsilon$. Für die endliche Menge $J_\epsilon = \{1, \dots, N_\epsilon\}$ gilt also $\sum_{k \in J_\epsilon} |a_k| > \sum_{k=1}^{\infty} |a_k| - \epsilon$, d.h. $\sum |a_k| - \epsilon$ ist keine obere Schranke von $\mathcal{F}$. $\square$

§ 16. REIHEN

Beweis von 16.12. Die endlichen Teilsummen der umgeordneten Reihe sind genau die gleichen wie die der ursprünglichen. Daher folgt die absolute Konvergenz der umgeordneten Reihe direkt aus Lemma 16.13. Wegen Proposition 16.8 existieren also sowohl

$$S := \sum_{k=1}^{\infty} a_k \quad \text{als auch} \quad T := \sum_{l=1}^{\infty} a_{k_l} \quad (16.21)$$

Für $\epsilon > 0$ sei nun N_ϵ so gross, dass

$$\left| S - \sum_{k=1}^{N_\epsilon} a_k \right| \leq \sum_{k=N_\epsilon+1}^{\infty} |a_k| < \epsilon \quad (16.22)$$

(Die Existenz von N_ϵ folgt aus der absoluten Konvergenz, die erste Abschätzung durch Anwenden von Lemma 16.6 auf die “Restreihe” $\sum_{k=N_\epsilon+1}^{\infty} a_k$.) Das Urbild von $\{1, \dots, N_\epsilon\}$ unter der Umordnung $l \mapsto k_l$ ist endlich, daher existiert ein M_ϵ so dass

$$\{k_1, \dots, k_{M_\epsilon}\} \supset \{1, \dots, N_\epsilon\} \quad (16.23)$$

Für alle $n \geq N_\epsilon$ und $m \geq M_\epsilon$ gilt dann

$$\left| \sum_{k=1}^n a_k - \sum_{l=1}^m a_{k_l} \right| \leq \sum_{k=N_\epsilon+1}^{\infty} |a_k| \quad (16.24)$$

Denn wegen (16.23) ist die Differenz der Partialsummen eine endliche Teilsumme der Restreihe $\sum_{k=N_\epsilon+1}^{\infty} a_k$.

• Da die Partialsummen in (16.24) für $n \rightarrow \infty$ bzw. $k \rightarrow \infty$ gegen S bzw. T konvergieren, folgt mit (16.22), dass $|S - T| \leq \sum_{k=N_\epsilon+1}^{\infty} |a_k| < \epsilon$, $\forall \epsilon > 0$. $\square$

Der *Grosse Umordnungssatz* ist eine Verallgemeinerung von 16.11, der Beweis sehr ähnlich.

Theorem 16.14. Sei $\sum_{k=1}^{\infty} a_k$ eine absolut konvergente Reihe und $(J^{(i)})_{i=1,2,\dots}$ eine (endliche oder unendliche) Familie von (endlichen oder unendlichen) Teilmengen von $\mathbb{N}$ mit $J^{(i)} \cap J^{(j)} = \emptyset$ für $i \neq j$ und $\cup_i J^{(i)} = \mathbb{N}$. Dann gilt

(i) Für jede Durchnummerierung (Abzählung/Anordnung) $(k_l^{(i)})_{l=1,2,\dots}$ der Elemente von $J^{(i)}$ ist $\sum_l a_{k_l^{(i)}}$ (entweder endlich oder) absolut konvergent, und ihr Wert $S^{(i)}$ unabhängig von der Durchnummerierung, $\forall i$.

(ii) Die (endliche Summe oder) Reihe $\sum_{i=1,2,\dots} S^{(i)}$ konvergiert absolut und ihr Wert ist gleich $\sum_{k=1}^{\infty} a_k =: S$.

(In der Rechenregel Lemma 16.3 kann $\sum a_k + \sum b_k$ als eine solche Umgruppierung der Reihe $\sum c_l$ mit $c_{2l-1} = a_l$, $c_{2l} = b_l$ aufgefasst werden. Die Umgruppierung ist erlaubt, wenn $\sum a_k$ und $\sum b_k$ getrennt konvergieren, absolute Konvergenz ist nicht notwendig. Die Umkehrung gilt natürlich nicht, d.h. $\sum (a_k + b_k)$ kann konvergieren, auch wenn die Reihen getrennt dies nicht tun. Ist beispielsweise $\forall k: a_k = 1$ und $b_k = -1$, so ist $\sum_{k=0}^{\infty} (a_k + b_k) = 0$, während die Reihen getrennt natürlich divergieren. Auch führt hier bereits eine kleine Verschiebung der beiden Reihen gegeneinander zur Konvergenz gegen einen anderen Wert: $a_0 + \sum_{k=1}^{\infty} (a_k + b_{k-1}) = 1$.)

Beweis von 16.14. (Beachte, dass die $J^{(i)}$ höchstens abzählbar unendlich sind und es höchstens abzählbar unendlich viele davon gibt, s. Def. 4.10.)

(i) Für endliches $J^{(i)}$ sind die Aussagen trivial. Andernfalls folgt die absolute Konvergenz wie eben aus dem Lemma 16.13, und die Unabhängigkeit von der Durchnummerierung aus Proposition 16.12. Wir fixieren dann für jedes i eine solche.

(ii) Für $\epsilon > 0$ sei N_ϵ wieder so, dass $\sum_{k=N_\epsilon+1}^\infty |a_k| < \epsilon$, so dass also wieder $|S - \sum_{k \in J} a_k| < \epsilon$ für jede endliche Teilmenge J von $\mathbb{N}$ mit $\{1, \dots, N_\epsilon\} \subset J$.

· Das Urbild von $\{1, \dots, N_\epsilon\}$ unter $(i, l) \mapsto k_l^{(i)}$ ist wieder beschränkt, d.h. es existieren H_ϵ und M_ϵ so gross, dass $\cup_{i=1}^{H_\epsilon} \{k_l^{(i)} \mid l \leq M_\epsilon\} \supset \{1, \dots, N_\epsilon\}$. Dann ist $\forall n \geq N_\epsilon, h \geq H_\epsilon, m \geq M_\epsilon$:

$$\left| \sum_{k=1}^n a_k - \sum_{i=1}^h \sum_{l \leq m} a_{k_l^{(i)}} \right| \leq \sum_{k=N_\epsilon+1}^\infty |a_k| < \epsilon \quad (16.25)$$

wobei wir vereinbaren, dass $a_{k_l^{(i)}} = 0$ falls $J^{(i)}$ endlich ist und $l > |J^{(i)}|$.

· Mit $m \rightarrow \infty$ und $n \rightarrow \infty$ folgt daraus wegen der Konvergenz der $\sum_l a_{k_l^{(i)}}$ für $i = 1, \dots, h$

$$\left| S - \sum_{i=1}^h S^{(i)} \right| < \epsilon \quad (16.26)$$

Dies gilt für alle $h \geq H_\epsilon, \forall \epsilon > 0$, so dass $S = \sum_i S^{(i)}$. (Gegebenenfalls kann diese Summe natürlich auch wieder endlich sein.) Die Konvergenz ist absolut, da $\sum_{i=1}^h \sum_{l \leq m} |a_{k_l^{(i)}}| \leq \sum_{k=1}^\infty |a_k|$ im Grenzübergang $m \rightarrow \infty$ $\sum_{i=1}^h |S^{(i)}| \leq \sum_{k=1}^\infty |a_k|$ impliziert. $\square$

· Ein beliebtes Anwendungsbeispiel von Theorem 16.14 ist die Summation von “Doppelreihen” der Form $(a_{kl})_{(k,l) \in \mathbb{N} \times \mathbb{N}}$: Das Resultat hängt nicht von der Abzählung $\mathbb{N} \xrightarrow{\cong} \mathbb{N} \times \mathbb{N}$ ab, sofern $\mathcal{F} = \{ \sum_{(k,l) \in J} |a_{kl}| \mid J \subset \mathbb{N} \times \mathbb{N} \text{ endlich} \}$ nach oben beschränkt ist.

· Konkret: $a_{kl} = \frac{1}{k^l}, k, l \geq 2$. Dann ist $\forall K, L$:

$$\sum_{k=2}^K \sum_{l=2}^L a_{kl} \leq \sum_{k=2}^K \sum_{l=2}^\infty \frac{1}{k^l} = \sum_{k=2}^K \frac{1}{k^2} \frac{1}{1 - \frac{1}{k}} \leq \sum_{k=2}^\infty \frac{1}{k(k-1)} \leq \sum_{k=1}^\infty \frac{1}{k^2} < \infty \quad (16.27)$$

(vgl. Beispiel 16.5) sodass beliebiges Vertauschen erlaubt ist. Es folgt

$$\sum_{l=2}^\infty (\zeta(l) - 1) = \sum_{l=2}^\infty \sum_{k=2}^\infty \frac{1}{k^l} = \sum_{k=2}^\infty \sum_{l=2}^\infty \frac{1}{k^l} = \sum_{k=2}^\infty \frac{1}{k(k-1)} = 1 \quad (16.28)$$

· Beachte: Hier wird bei der Berechnung der teleskopierenden Reihe

$$\sum_{k=1}^\infty \frac{1}{k(k+1)} = \sum_{k=1}^\infty \left(\frac{1}{k} - \frac{1}{k+1} \right) \quad (16.29)$$

die Umordnung nicht über 16.14 begründet (die Teilreihen konvergieren ja auch nicht), sondern durch eine explizite Berechnung der Partialsummen:

$$\sum_{k=1}^n \left(\frac{1}{k} - \frac{1}{k+1} \right) = \sum_{k=1}^n \frac{1}{k} - \sum_{k=2}^{n+1} \frac{1}{k} = 1 - \frac{1}{n+2} \xrightarrow{n \rightarrow \infty} 1 \quad (16.30)$$

Konvergenzkriterien

Wir geben nun zwei einfache Kriterien an, die es erlauben bei konkret gegebenen Reihen mit fester oder geschickt gewählter Reihenfolge der zu summierenden Reihenglieder zwischen Konvergenz und Divergenz zu entscheiden. — Wie oben gezeigt ist es zwar notwendig, aber nicht hinreichend (vgl. (16.10) mit Lemma 16.2) für die Konvergenz, dass die Reihenglieder eine Nullfolge bilden. (Diese Eigenschaft ist unabhängig von der gewählten Reihenfolge.) Der Vergleich mit der geometrischen Reihe Beispiel 16.4 suggeriert jedoch, dass es für die (absolute) Konvergenz einer Reihe $\sum a_k$ ausreicht, wenn die Beträge $|a_k|$ der Reihenglieder für grosse k mindestens so schnell gegen Null gehen wie die Potenzen L^k einer Zahl $L < 1$. Zur Gewinnung von L aus den Reihengliedern können wir entweder $|a_{k+1}/a_k|$ betrachten (Quotientenkriterium), oder $|a_k|^{1/k}$ (Wurzelkriterium). Zur Quantifizierung benutzen wir den Begriff des limes superior aus Def. 14.18 und seine verschiedenen Charakterisierungen aus Prop. 14.19, die wir hier kurz wiederholen.

-
- Ist $(x_k) \subset X$ eine Folge in einem metrischen Raum, so heisst $y \in X$ Häufungswert der Folge, falls eine Teilfolge $(x_{k_l})_l$ von (x_k) existiert, die gegen y konvergiert. Es gilt: $y \in X$ ist genau dann ein Häufungswert, wenn für alle $\epsilon > 0$ $x_k \in B_\epsilon(x)$ für unendlich viele k .
 - Ist $(x_k) \subset \mathbb{R}$ eine *beschränkte* Folge *reeller* Zahlen, so existiert nach Bolzano-Weierstrass 14.17 mindestens ein Häufungswert. Ausserdem ist die Menge aller Häufungswerte beschränkt. Man nennt das Supremum der Häufungswerte den Limes Superior von (x_k) .

$$\begin{aligned} \limsup x_k &= \sup\{y \in \mathbb{R} \mid \text{Es existiert eine gegen } y \text{ konvergente Teilfolge}\} \\ &= \max\{y \in \mathbb{R} \mid \text{Es existiert eine gegen } y \text{ konvergente Teilfolge}\} \end{aligned} \quad (16.31)$$

Beh.: Das Supremum wird angenommen, d.h. es existiert eine Teilfolge (x_{k_m}) welche gegen $z := \limsup x_k$ konvergiert.

Bew.: Für jedes $m \in \mathbb{N}$ existiert ein Häufungswert y_m mit $y_m > z - \frac{1}{m}$ (Def. des Supremums) und es gibt unendlich viele x_k mit $|x_k - y_m| < \frac{1}{m}$ (Def. des Häufungswerts). Man findet dann sukzessive eine Folge k_m so, dass $|x_{k_m} - z| < \frac{2}{m}$ (Dreiecksungleichung) und $k_{m+1} > k_m \forall m$, und (x_{k_m}) ist eine gegen z konvergente Teilfolge. $\square$

Variante: Für jedes $z' > z = \limsup x_k$ ist $x_k > z'$ höchstens für endlich viele k .

Bew: Andernfalls existiert eine Teilfolge von (x_k) , deren Glieder alle grösser als z' sind. Jeder Häufungswert dieser Teilfolge wäre dann grösser oder gleich z' , entgegen der Definition von z als dem grössten Häufungswert. $\square$

- Für ein nach oben unbeschränkte Folge setzt man formal $\limsup = \infty$.
-

Proposition 16.15. *Es sei $(a_k)_{k \in \mathbb{N}}$ ein Folge komplexer Zahlen, und*

$$L := \limsup |a_k|^{1/k} \in \mathbb{R}_{\geq 0} \cup \{\infty\} \quad (16.32)$$

Dann gilt

- (i) *Für $L < 1$ konvergiert die Reihe $\sum a_k$ absolut.*
- (ii) *Für $L > 1$ divergiert die Reihe.*

Beweis. Für $L = \infty$ ist $(|a_k|^{1/k})$ unbeschränkt, und die Divergenz der Reihe klar.

- Für $L < 1$ existiert ein $q \in (L, 1)$ (für $L > 0$ z.B. $q := \sqrt{L}$). Nach der Variante oben ist dann $|a_k|^{1/k} > q$ nur für endlich viele k , d.h. es existiert ein N so, dass

$|a_k|^{1/k} \leq q \ \forall k \geq N$. Dann ist

$$\sum_{k=N}^{\infty} q^k = \frac{q^N}{1-q} \quad (16.33)$$

eine konvergente Majorante für $\sum_{k=N}^{\infty} a_k$. Mit Lemma 16.6 folgt daraus (i).

· Falls $\infty > L > 1$, so gibt es $\epsilon > 0$ mit $L - \epsilon > 1$, und unendlich viele k mit $|a_k|^{1/k} \geq L - \epsilon > 1$. Die a_k sind dann nicht einmal eine Nullfolge, also kann die Reihe $\sum a_k$ nicht konvergieren. Dies impliziert (ii). $\square$

Beispiel 16.16. Sei $r \in \mathbb{Q}$ und $z \in \mathbb{C}$. Dann ist wegen $\lim_{k \rightarrow \infty} k^{1/k} = 1$ (vgl. Beispiel 14.8 (iii))

$$\limsup |k^r z^k|^{1/k} = \lim_{k \rightarrow \infty} k^{r/k} |z| = |z| \quad (16.34)$$

und daher konvergiert

$$\sum_{k=1}^{\infty} k^r z^k \quad (16.35)$$

für $|z| < 1$ absolut und divergiert für $|z| > 1$, genau wie die geometrische Reihe.

· Um (als Beispiel) für $r = 1$ die Reihe zu berechnen, schreiben wir

$$\sum_{k=1}^{\infty} k z^k = z \sum_{k=0}^{\infty} (k+1) z^k = z \sum_{k=0}^{\infty} \left(\sum_{l=0}^k z^k \right) \quad (16.36)$$

als eine unendliche Summe über $k = 0, 1, \dots$ der endlichen Summen $(k+1)z^k = \sum_{l=0}^k z^k$. Wegen der absoluten Konvergenz können wir über diese Indexmenge auch in einer beliebigen anderen Parametrisierung summieren, ohne das Ergebnis zu ändern (s. Thm. 16.14). Unter der Bijektion

$$\{(k, l) \mid (k, l) \in \mathbb{N}_0 \times \mathbb{N}_0, l \leq k\} \ni (k, l) \mapsto (k-l, l) =: (m, l) \in \mathbb{N}_0 \times \mathbb{N}_0 \quad (16.37)$$

(Umkehrabbildung: $k = m + l$) wird aus (16.36) insbesondere

$$\begin{aligned} z \sum_{l \leq k} z^k &= z \sum_{(m, l) \in \mathbb{N}_0 \times \mathbb{N}_0} z^{l+m} = z \sum_{l=0}^{\infty} z^l \left(\sum_{m=0}^{\infty} z^m \right) \\ &= z \left(\sum_{l=0}^{\infty} z^l \right)^2 = \frac{z}{(1-z)^2} \end{aligned} \quad (16.38)$$

Allgemeiner ist für jedes Polynom $p(k)$ die Reihe $\sum p(k)z^k$ für $|z| < 1$ konvergent und eine rationale Funktion von z . (Übungsaufgabe, bzw. gliedweises Ableiten im Anschluss an ??.)

Korollar 16.17. *Es existiere $\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| =: q$. Dann gilt*

(i) *Für $q < 1$ konvergiert $\sum a_k$ absolut.*

(ii) *Für $q > 1$ divergiert die Reihe.*

§ 16. REIHEN

Beweis. (Versteckt in der Voraussetzung ist die Aussage $a_k \neq 0$ für fast alle k ; $q = 0$ ist hingegen erlaubt.)

(i) Wir zeigen $\limsup |a_k|^{1/k} \leq q$ (und verweisen dann auf das Wurzelkriterium). Für gegebenes $q' > q$ sei K so, dass

$$\left| \frac{a_{k+1}}{a_k} \right| < q' \quad \text{für alle } k \geq K \quad (16.39)$$

Dann gilt für solche k

$$\begin{aligned} |a_k| &= \left| \frac{a_k}{a_{k-1}} \right| \cdots \left| \frac{a_{K+1}}{a_K} \right| \cdot |a_K| \leq (q')^{k-K} \cdot |a_K| \\ \Rightarrow |a_k|^{1/k} &\leq q' \cdot (q'^{-K} |a_K|)^{1/k} \end{aligned} \quad (16.40)$$

Wegen 14.8(ii) folgt daraus $\limsup |a_k|^{1/k} \leq q'$. Da dies für alle $q' > q$ gilt, folgt $\limsup |a_k|^{1/k} \leq q$.

(ii) Falls $q > 1$, so wachsen die $|a_k|$ ab einem bestimmten Index streng monoton, die a_k bilden also nicht einmal eine Nullfolge. $\square$

Beispiel 16.18. · Die Exponentialreihe $\sum_{k=0}^{\infty} \frac{z^k}{k!}$ konvergiert (für $z = 0$ sowieso und sonst) wegen

$$\frac{z^{k+1}}{(k+1)!} \cdot \frac{k!}{z^k} = \frac{z}{k+1} \xrightarrow{k \rightarrow \infty} 0 \quad (16.41)$$

für alle $z \in \mathbb{C}$ absolut.

· Die Konvergenz der hypergeometrischen Reihe

$$F(a, b, c; z) := \sum_{k=0}^{\infty} \frac{(a)_k (b)_k}{(c)_k k!} z^k \quad (16.42)$$

wird in den Übungen diskutiert.

· Man beachte, dass die Kriterien 16.15 und 16.17 für $L = 1$ bzw. $q = 1$ nicht schlüssig sind. Es gibt eine Fülle an Beispielen und eine Reihe feinere Kriterien zur Entscheidung zwischen Konvergenz und Divergenz auch in solchen Fällen.

· Das Wurzelkriterium ist echt stärker als das Quotientenkriterium. Selbst im Fall $\lim |a_k|^{1/k} = L < 1$ folgt daraus nicht, dass $\lim |a_{k+1}/a_k|$ existiert oder < 1 . Beispiel: $1 + L^3 + L^2 + L^5 + L^4 + L^7 + \cdots$ für $L < 1$.

Potenzreihen

Das Interesse der meisten obigen Beispielen, speziell 16.4, (16.36), (16.41), (16.42) entsteht durch die Abhängigkeit der Reihenglieder von der komplexen Zahl, z . Konsequenterweise fasst man Reihen der Form

$$P(z) = \sum_{n=0}^{\infty} a_n z^n \quad (16.43)$$

als unendliche Linearkombinationen der elementaren Potenzfunktionen $z \mapsto z^n$ (zunächst: $a_n, z \in \mathbb{C}$) auf, m.a.W. als Reihen im Vektorraum der komplexwertigen Funktionen auf $\mathbb{C}$.

· Wie wir gleich sehen werden, ist die Dichotomie zwischen Konvergenz und Divergenz in Abhängigkeit von $|z|$ eine ganz allgemeine Eigenschaft von solchen Potenzreihen. Auf den Teilmengen von $\mathbb{C} \ni z$, auf denen sie konvergieren (intuitiv: “kleine $|z|$ ”) können Potenzreihen zur Definition einer grossen Klasse von nicht elementaren (“transzendenten”) Funktionen benutzt werden. Die Bedeutung solcher Reihenentwicklungen für die theoretische Physik kann schwerlich überschätzt werden.

Lemma 16.19. *Angenommen, $P(z)$ konvergiert in einem Punkt $z_0 \in \mathbb{C}$ mit $z_0 \neq 0$. Dann konvergiert $P(z)$ in jedem Punkt $z \in \mathbb{C}$ mit $|z| < |z_0|$ absolut.*

Beweis. Die Reihenglieder $a_n z_0^n$ bilden eine Nullfolge, es existiert also insbesondere ein S mit $|a_n z_0^n| \leq S$ für alle n . Daraus folgt für alle $|z| < |z_0|$ mit $q = \left| \frac{z}{z_0} \right| < 1$:

$$|a_n z^n| = |a_n z_0^n|^n \cdot \left| \frac{z}{z_0} \right|^n \leq S q^n \quad (16.44)$$

$P(z)$ besitzt also die konvergente Majorante $\sum S q^n = \frac{S}{1-q}$ und konvergiert daher absolut. $\square$

Definition 16.20. Für gegebene Potenzreihe $P(z) = \sum a_n z^n$ setzen wir

$$R = \sup \{ r \in \mathbb{R} \mid P(r) \text{ konvergiert} \} \quad (16.45)$$

(Für $r = 0$ konvergiert die Reihe trivialerweise. Daher ist $R \geq 0$. Ist die Menge rechts unbeschränkt, so setzen wir formal $R = \infty$.) Wir nennen R den *Konvergenzradius* und $B_R(0) = \{z \mid |z| < R\}$ die *Konvergenzscheibe*, denn es gilt:

Theorem 16.21. *Für $|z| < R$ konvergiert $P(z)$ absolut.*

· *Für $|z| > R$ divergiert $P(z)$.*

Beweis. Für $|z| < R$ existiert ein $r' > 0$ mit $|z| < r' < R$. Da $P(r')$ für jedes $0 < r' < R$ konvergiert, impliziert 16.19 die absolute Konvergenz von $P(z)$.

· Wäre für ein $|z| > R$ die Reihe $P(z)$ konvergent, so wäre wegen Lemma 16.19 für jedes $R' < |z|$ die Reihe $P(R')$ ebenfalls konvergent (sogar absolut). Wählt man ein solches R' mit $|z| > R' > R$, so erhält man einen Widerspruch zur Supremumsdefinition von R . $\square$

Wurzel- und Quotientenkriterium 16.15, 16.17 geben die folgenden essentiellen Formeln zur Berechnung des Konvergenzradius R einer Potenzreihe $\sum a_n z^n$:

$$\begin{aligned} R &= \frac{1}{\limsup |a_n|^{1/n}} && \text{(Formel von Cauchy-Hadamard)} \\ R &= \frac{1}{\lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right|} && \text{falls der Grenzwert existiert} \end{aligned} \quad (16.46)$$

In diesen Formeln gilt die Vereinbarung, dass $\frac{1}{0} = \infty$ und $\frac{1}{\infty} = 0$.

§ 16. REIHEN

Beispiel. $\sum \frac{z^n}{n!}$ hat $R = \infty$, konvergiert also überall. $\sum n!z^n$ hat $R = 0$ und konvergiert nur für $z = 0$.

In der Praxis ist insbesondere auch die Güte der Approximation einer Potenzreihe durch ihre Partialsummen von Interesse.

Lemma 16.22. Für jedes $N \in \mathbb{N}$ und $r < R$ existiert eine Konstante $C \geq 0$ so, dass für alle $|z| \leq r$

$$\left| \sum_{n=0}^{\infty} a_n z^n - \sum_{n=0}^N a_n z^n \right| \leq C \cdot |z|^{N+1} \quad (16.47)$$

In Worten: Die Differenz zwischen der vollen Potenzreihe und einer endlichen Teilsumme kann bis auf eine von z unabhängige multiplikative Konstante durch den ersten weggelassenen Term abgeschätzt werden. Dies ist besonders nützlich für $|z|^{N+1} \ll C^{-1}$.

Beweis von Lemma 16.22. Wähle ρ mit $r < \rho < R$. Da $\sum a_n \rho^n$ konvergiert, existiert eine Konstante S so, dass $|a_n \rho^n| \leq S \forall n$. Für $|z| \leq r$ konvergiert die Potenzreihe und daher auch die Restreihe $\sum_{n=N+1}^{\infty} a_n z^n$ absolut und es gilt die Abschätzung

$$\begin{aligned} \left| \sum_{n=N+1}^{\infty} a_n z^n \right| &\leq \sum_{n=N+1}^{\infty} |a_n z^n| = \sum_{n=N+1}^{\infty} |a_n \rho^n| \frac{|z|^n}{\rho^n} \\ &\leq S \cdot \frac{|z|^{N+1}}{\rho^{N+1}} \sum_{n=0}^{\infty} \frac{r^n}{\rho^n} = |z|^{N+1} \cdot \frac{S}{\rho^{N+1}(1 - r/\rho)} \end{aligned} \quad (16.48)$$

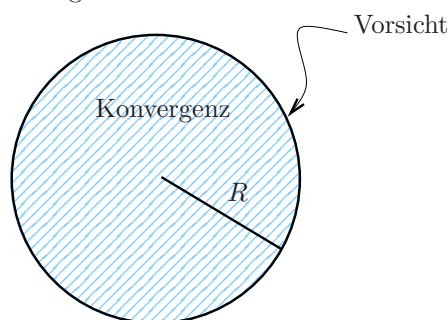
Daraus folgt die Behauptung. □

Merkwissen

I. Unter absoluter (dominierter/majorisierter) Konvergenz sind alle vernünftigen Manipulationen erlaubt, im Allgemeinen aber muss man vorsichtig sein.

II. Jede Potenzreihe kommt mit einem Konvergenzradius (der auch 0 oder ∞ sein kann). Innerhalb der Konvergenzscheibe konvergiert eine Potenzreihe absolut, und der Fehler lässt sich durch den ersten weggelassenen Term abschätzen. Auf der Konvergenzkreislinie ist wieder Vorsicht geboten.

Divergenz



III. Die Definitionen und Konvergenz-Kriterien übertragen sich mit notwendigen Änderungen auf Reihen in einem vollständigen normierten Vektorraum $(V, \|\cdot\|)$ über $\mathbb{R}$ oder $\mathbb{C}$. In diesem Zusammenhang heisst etwa eine Reihe $\sum v_k$ mit $v_k \in V$ absolut konvergent (man sagt auch konvergent in der Norm), wenn die Reihe der Normen $\sum \|v_k\|$ konvergiert. Insbesondere gelten Lemma 16.6 und die Umordnungsregeln sinngemäss. Das Leibniz-Kriterium 16.10 gilt natürlich nur in $\mathbb{R}$!

III'. Häufig anzutreffen sind Potenzreihen, in denen die Veränderliche einer nicht notwendig kommutativen Banach-Algebra $(\mathcal{A}, \|\cdot\|)$ entstammt (z.B. $\text{Mat}_{n \times n}(\mathbb{R})$), während die Koeffizienten im Grundkörper bleiben.⁹¹ Definiert man R wie in (16.46), so konvergiert aus den gleichen Gründen wie oben eine solche Potenzreihe innerhalb von $B_R(0)$ in der Norm. Da die Norm aber im allgemeinen nicht multiplikativ, sondern nur submultiplikativ ist, kann man nicht auf Divergenz überall ausserhalb von $B_R(0)$ schliessen. Beispiel: Die Matrix $A = \begin{pmatrix} 0 & C \\ 0 & 0 \end{pmatrix}$ liegt für C gross genug ausserhalb von $B_1(0)$, unabhängig von der gewählten Norm auf $\text{Mat}_{2 \times 2}(\mathbb{R})$. Wegen $A^2 = 0$ konvergiert aber die geometrische Reihe $\sum A^k = \begin{pmatrix} 1 & C \\ 0 & 1 \end{pmatrix} = (1 - A)^{-1}$ für jedes C .

Beispiele

- Geometrische (oder Neumann-)Reihe

$$\sum_{n=0}^{\infty} x^n = (1 - x)^{-1} \quad \text{für } \|x\| < 1 \quad (16.49)$$

- Logarithmus-(oder Mercator-)Reihe

$$\sum_{n=1}^{\infty} (-1)^{n-1} \frac{x^n}{n} \quad R = 1 \quad (16.50)$$

- Exponentialreihe

$$\exp(x) := \sum_{n=0}^{\infty} \frac{x^n}{n!} \quad R = \infty \quad (16.51)$$

- Sinus-Reihe

$$\sin(x) := \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n+1}}{(2n+1)!} \quad R = \infty \quad (16.52)$$

- Cosinus-Reihe

$$\cos(x) := \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n}}{(2n)!} \quad R = \infty \quad (16.53)$$

- Euler-Formel (für Grundkörper $= \mathbb{C}$)

$$\exp(ix) = \cos(x) + i \sin(x) \quad (16.54)$$

Proposition 16.23. Für alle $x, y \in \mathcal{A}$ mit $xy = yx$ gilt

$$\exp(x) \exp(y) = \exp(x + y) \quad (16.55)$$

und daraus abgeleitet die bekannten Additionstheoreme der via (16.54) definierten trigonometrischen Funktionen.

⁹¹Wir nehmen die Existenz eines Einselements $1 \in \mathcal{A}$ an, sodass $x^0 = 1$ erklärt ist.

§ 16. REIHEN

Lemma 16.24. Sei $(x_n) \subset \mathcal{A}$ eine Folge mit $\lim_{n \rightarrow \infty} x_n = x$. Dann gilt

$$\lim_{n \rightarrow \infty} \left(1 + \frac{x_n}{n}\right)^n = \exp(x) \quad (16.56)$$

Beweis von 16.23. (unter der Annahme, dass 16.24 gilt)

$$\begin{aligned} \exp(x) \exp(y) &= \lim_{n \rightarrow \infty} \left(1 + \frac{x}{n}\right)^n \lim_{n \rightarrow \infty} \left(1 + \frac{y}{n}\right)^n = (\text{Rechenregeln für Folgen}) \\ &= \lim_{n \rightarrow \infty} \left(1 + \frac{x}{n}\right)^n \left(1 + \frac{y}{n}\right)^n \quad (\text{wegen } xy = yx) \\ &= \lim_{n \rightarrow \infty} \left(1 + \frac{x}{n} + \frac{y}{n} + \frac{xy}{n^2}\right)^n \end{aligned} \quad (16.57)$$

Wegen $\lim_{n \rightarrow \infty} \left(x + y + \frac{xy}{n}\right) = x + y$ folgt die Behauptung. $\square$

Beweis von 16.24. Für $\epsilon > 0$ sei N so gross, dass

$$\sum_{k=N+1}^{\infty} \frac{(\|x\| + 1)^k}{k!} < \epsilon \quad \text{und} \quad \|x_n\| < \|x\| + 1 \quad \forall n > N \quad (16.58)$$

Dann gilt für $n \geq N$:

$$\begin{aligned} &\left\| \left(1 + \frac{x_n}{n}\right)^n - \exp(x) \right\| \\ &\leq \underbrace{\sum_{k=0}^N \left\| \binom{n}{k} \frac{x_n^k}{n^k} - \frac{x^k}{k!} \right\|}_{\rightarrow 0 \text{ für } n \rightarrow \infty} + \underbrace{\sum_{k=N+1}^n \left\| \binom{n}{k} \frac{x_n^k}{n^k} \right\|}_{\leq \frac{(\|x\|+1)^k}{k!} < \epsilon} + \underbrace{\sum_{k=N+1}^{\infty} \frac{\|x^k\|}{k!}}_{\leq \frac{(\|x\|+1)^k}{k!} < \epsilon} \end{aligned} \quad (16.59)$$

wobei wir benutzt haben, dass

$$\binom{n}{k} \frac{1}{n^k} = \frac{1}{k!} \prod_{i=0}^{k-1} \left(1 - \frac{i}{n}\right) \begin{cases} \leq \frac{1}{k!} & \text{für den mittleren Term} \\ \xrightarrow{n \rightarrow \infty} \frac{1}{k!} & \text{für den ersten Term} \end{cases} \quad (16.60)$$

Es gibt also ein $\tilde{N} > N$ so, dass für $n \geq \tilde{N}$ der erste Term auf der rechten Seite von (16.59) $< \epsilon$ ist, alles zusammen $< 3\epsilon$. $\square$

· Insbesondere ist $\exp(x) \exp(-x) = \exp(0) = 1$ und daher

$$\exp(x) \neq 0 \quad \forall x \quad (16.61)$$