



UNIVERSITÄT
HEIDELBERG
ZUKUNFT
SEIT 1386

SKRIPT ZUR VORLESUNG
HÖHERE MATHEMATIK II
FÜR DAS
STUDIUM DER PHYSIK

gehalten an der
Universität Heidelberg
im
Sommersemester 2020
von
Johannes Walcher

INHALTSVERZEICHNIS

1	EINFÜHRUNG	3
§ 1	Die reellen Zahlen	3
§ 2	Die komplexen Zahlen	12
2	KONVERGENZ	17
§ 3	Folgen reeller Zahlen	17
§ 4	Metrische Räume	27
§ 5	Reihen	43
3	STETIGKEIT	59
§ 6	Stetige Abbildungen	60
§ 7	Werte stetiger Funktionen	65
§ 8	Gleichmässigkeit	70
4	DIFFERENTIATION	75
§ 9	Eine reelle Veränderliche	75
§ 10	Mehrere Veränderliche	80
§ 11	Komplex differenzierbare Funktionen	93
5	UMKEHRUNG	95
§ 12	Determinante und Spur	95
§ 13	Gleichungen im $\mathbb{R}^n$	98
§ 14	Stammfunktionen	107
§ 15	Gewöhnliche Differentialgleichungen	116

KAPITEL 1

EINFÜHRUNG

Diese Vorlesung richtet sich an Studierende der Physik im zweiten Semester. Das Hauptziel ist die Bereitstellung der für das weitere Studium der Physik wichtigsten Konzepte und Rechenmethoden der *Mathematischen Analysis*. Dabei baut die Vorlesung einerseits auf den im ersten Semester vermittelten Grundlagen der linearen Algebra auf (insbesondere werden die üblichen Grundgedanken zur mathematischen Beweisführung und Notation und zur elementaren Mengenlehre *nicht* im Detail besprochen), und stellt andererseits einen Bezug zu den theoretischen Physik-Vorlesungen her, indem einige der dort bereits verwendeten Begriffe und Ergebnisse hier zunächst noch einmal neu begründet, und anschliessend weiterentwickelt werden. (Dabei soll ganz und gar nicht die gewohnte Anschauung im Keime erstickt werden, sondern vielmehr die Wertschätzung und das Vertrauen in die Methoden der “anderen” Disziplin gestärkt werden.)

Für Kommentare, auch Fehlermeldungen, zum Skript, per Email an walcher@uni-heidelberg.de, bin ich sehr dankbar. Man beachte aber, dass dieses Skript keinen Anspruch auf Abgeschlossenheit erhebt: Viele Beweise werden bestenfalls angedeutet, mancher Begriff nur indirekt definiert.

· Zur Homepage der Vorlesung: www.mathi.uni-heidelberg.de/~walcher/teaching/sose20/hoemaphys/

§ 1 Die reellen Zahlen

Die reellen Zahlen bilden die Grundlage der Analysis, und damit der gesamten theoretischen Physik. Sie werden durch drei Gruppen von Eigenschaften charakterisiert.

Arithmetik

Definition 1.1. Ein *Körper* ist eine Menge $\mathbb{K}$ mit zwei ausgezeichneten und verschiedenen Elementen $0 \in \mathbb{K}$ und $1 \in \mathbb{K}$, $0 \neq 1$, zusammen mit zwei binären Verknüpfungen, welche wir auch “Rechenoperationen” nennen und wie gewöhnlich notieren:

$$\begin{aligned} \text{Addition: } \mathbb{K} \times \mathbb{K} \ni (a, b) &\mapsto a + b \in \mathbb{K} \\ \text{Multiplikation: } \mathbb{K} \times \mathbb{K} \ni (a, b) &\mapsto a \cdot b := ab \in \mathbb{K} \end{aligned} \tag{1.1}$$

und welche die folgenden Rechenregeln erfüllen:

$$\begin{aligned}
 \text{Kommutativität: } & \forall a, b \in \mathbb{K} : a + b = b + a, \quad a \cdot b = b \cdot a \\
 \text{Assoziativität: } & \forall a, b, c \in \mathbb{K} : (a + b) + c = a + (b + c), \quad (a \cdot b) \cdot c = a \cdot (b \cdot c) \\
 \text{Neutrale Elemente: } & \forall a \in \mathbb{K} : a + 0 = a, \quad a \cdot 1 = a \\
 \text{Distributivgesetz: } & \forall a, b, c \in \mathbb{K} : a \cdot (b + c) = a \cdot b + a \cdot c
 \end{aligned} \tag{1.2}$$

derart, dass die Gleichungen

$$a + x = 0, \quad b \cdot y = 1 \tag{1.3}$$

für alle $a \in \mathbb{K}$ und $b \in \mathbb{K} \setminus \{0\}$ Lösungen $x, y \in \mathbb{K}$ besitzen.

Lemma 1.2. (i) Die Lösungen der Gleichungen (1.3) sind eindeutig. Wir bezeichnen sie mit $-a$ bzw. $b^{-1} =: \frac{1}{b}$, und schreiben natürlich auch $b - a := b + (-a)$ etc..

(ii) Insbesondere sind 0 und 1 eindeutig.

(iii) Ausserdem impliziert $xy = 0$, dass mindestens $x = 0$ oder $y = 0$.

Beweis. (i) Aus $a + x = 0$ und $a + x' = 0$ folgt unter Anwendung nur der aufgeführten Regeln (1.2):

$$x = x + 0 = x + (a + x') = (x + a) + x' = (a + x) + x' = 0 + x' = x' \tag{1.4}$$

· Ein ähnliches Argument zeigt die Eindeutigkeit der Lösung von $by = 1$. (Die Existenz einer Lösung für $b = 0$ stünde im Widerspruch zu $0 \cdot y = 0 \neq 1 \forall y \in \mathbb{K}$.)

(ii) Aus $a + 0' = a \forall a$ folgt $0' = 0 + 0' = 0$ und ähnlich für die 1.

(iii) Ist $xy = 0$ und $x \neq 0$ so folgt $0 = x^{-1}(xy) = (x^{-1}x)y = 1 \cdot y = y$. $\square$

Bemerkungen. · Die Motivation zur Forderung dieser Gesetze der Arithmetik sind zunächst einmal die Eigenschaften der natürlichen Zahlen $\mathbb{N} = \{1, 2, 3, \dots\}$, in denen man Assoziativität, Kommutativität, die Eigenschaft der 1, und Distributivität axiomatisch begründen kann, nach Geschmack auch mit der 0.

· Die Erweiterung von den natürlichen Zahlen $\mathbb{N}$ zu den ganzen Zahlen $\mathbb{Z} = \{0, 1, -1, 2, -2, 3, -3, \dots\}$, macht die erste der Gleichungen (1.3) lösbar, der Übergang zu den rationalen Zahlen $\mathbb{Q} = \{\frac{p}{q} \mid p \in \mathbb{Z}, q \in \mathbb{N}, (p, q) = 1\}$ dann die zweite. Es ist bemerkenswert, dass dieser Prozess an dieser Stelle abgeschlossen werden kann: $\mathbb{Q}$ ist ein Körper und muss dazu nicht etwa noch weitererweitert werden.

· Es wäre dies jedoch nicht die einzige Möglichkeit. Anstatt zu erweitern kann man auch zu Restklassen übergehen: Für jede Primzahl p ist $\mathbb{F}_p := \mathbb{Z}/(p\mathbb{Z})$ ein endlicher Körper.

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

· In den Naturwissenschaften dienen Körper zur abstrakten Buchhaltung konkret ausgeführter endlicher Messoperationen. Genauer gesagt enthalten die Bildmengen physikalischer Theorien in der Regel¹ Vektorräume, in denen physikalische Grössen

¹und in einer gewissen Näherung. Der Leitgedanke dieser Vorlesung ist es, den Begriff der Näherung mathematisch genauer zu fassen.

§1. DIE REELLEN ZAHLEN

mit Zahlen multipliziert und untereinander addiert werden können, während die Multiplikation zweier physikalischer Größen miteinander von anderer Natur ist. (Das Produkt zweier Längen ist z.B. gar keine Länge!) Man vergleiche dies mit dem operationellen Erlernen der Grundrechenarten.

· Eine wichtige Eigenschaft von Messprozessen ist die quantitative Vergleichbarkeit physikalischer Größen. Im mathematischen Bild führt man diese Vergleiche (eventuell mit weiteren Zwischenschritten, siehe unten!) auf den folgenden Begriff zurück.

Anordnung

Definition 1.3. Ein *angeordneter Körper* ist ein Körper $\mathbb{K}$ zusammen mit einer Teilmenge P (den “positiven” Elementen) derart, dass

$$\begin{aligned}\mathbb{K} &= P \cup \{0\} \cup (-P) \\ P + P &\subset P \\ P \cdot P &\subset P\end{aligned}\tag{1.5}$$

Wir schreiben auch $a > 0$ oder $0 < a$ für $a \in P$ und $a < 0$ oder $0 > a$ für $a \in -P$ ($\Leftrightarrow -a \in P$).

Bemerkungen. · In dieser Definition bedeutet $\cup$ “disjunkte Vereinigung”, das heisst für drei Mengen A, B, C gilt $A = B \cup C \Leftrightarrow A = B \cup C$ und $B \cap C = \emptyset$.

· Die Definition der Anordnung durch Auszeichnung einer Teilmenge mag zunächst etwas seltsam erscheinen. Wie das folgende Lemma zeigt, garantiert die in (1.5) formulierte Verträglichkeit mit Addition und Multiplikation, dass man aus P die gewöhnlichen Vergleichsrelationen rekonstruieren kann. Formal sind Relationen, speziell auch Abbildungen und binäre Verknüpfungen ja ebenfalls durch Auszeichnung von Teilmengen charakterisiert.

Lemma 1.4. (i) Durch die Vorschrift $a > b : \Leftrightarrow a - b \in P$ wird auf einem angeordneten Körper eine strenge Totalordnung definiert (transitiv und trichotomisch), bzw. durch $a \geq b : \Leftrightarrow (a > b \text{ oder } a = b)$ eine Totalordnung (transitiv, antisymmetrisch und total) und es gelten die üblichen (zum Teil empfindlichen) Rechenregeln für Ungleichungen, wie z.B.

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

$$1 > 0\tag{1.6}$$

Beweis. (i) Die Trichotomie: $\forall a, b \in \mathbb{K}$ gilt entweder $a > b$ oder $a = b$ oder $b > a$ folgt aus der dreiteiligen Zerlegung in (1.5). Die Transitivität sehen wir so:

$$\begin{aligned}a > b \text{ und } b > c &\Leftrightarrow a - b \in P \text{ und } b - c \in P \Rightarrow \\ &(\text{wegen } P + P \subset P) \Rightarrow (a - b) + (b - c) = a - c \in P \Leftrightarrow a > c\end{aligned}\tag{1.7}$$

(ii), (iii) Übungsaufgaben

(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 $1 = (-1)^2 > 0$ $\square$

Bemerkungen. · Die Forderung nach Anordnung schliesst endliche Körper wie $\mathbb{F}_p$ aus: In der Tat enthält jeder angeordnete Körper die natürlichen Zahlen, d.h. die eindeutige strukturerhaltende Abbildung $\phi : \mathbb{N} \rightarrow \mathbb{K}$ (mit $\mathbb{N} \ni 1 \mapsto \phi(1) := 1 \in \mathbb{K}$, $\mathbb{N} \ni 2 \mapsto \phi(2) := 1 + 1 \in \mathbb{K}$ etc.) ist injektiv: Wegen (1.6) (in $\mathbb{K}$) und den Anordnungsaxiomen (1.5) gilt auch $\phi(n) > 0$, also $\phi(n) \neq 0 \forall n \in \mathbb{N}$ (vollständige Induktion). Wir schreiben natürlich $\phi(n) =: n \in \mathbb{K}$. Ebenso enthält jeder angeordnete Körper die ganzen Zahlen und die rationalen Zahlen als angeordneten Unterkörper. · Für Mathematiker wäre es aber, wie schon bei den Rechenregeln zuvor, wichtig zu betonen, dass nicht alle möglichen angeordneten Körper alle von den rationalen Zahlen her gewohnten Eigenschaften besitzen.

Ein einfaches Hilfsmittel:

Lemma 1.5. Für jedes $x \geq -1$ in einem angeordneten Körper $\mathbb{K}$ gilt für alle $n \in \mathbb{N}$:

$$(1 + x)^n \geq 1 + nx \quad (1.8)$$

Beweis. Auf der rechten Seite ist natürlich $(1 + x)^n$ 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:

$$(1 + x)^{n+1} \geq (1 + nx)(1 + x) = 1 + (n + 1)x + nx^2 \geq 1 + (n + 1)x \quad (1.9)$$

$\square$

Eine Sprachregelung:

Definition 1.6. Sei $\mathbb{K}$ ein angeordneter Körper. Eine Teilmenge $T \subset \mathbb{K}$ heisst *nach oben (unten) beschränkt*, falls eine *obere (untere) Schranke* von T existiert, d.h. ein $B \in \mathbb{K}$ s.d.

$$\forall x \in T : B \geq x \quad (B \leq x) \quad (1.10)$$

Wir sagen “beschränkt” für eine Teilmenge, die sowohl nach unten als auch nach oben beschränkt ist.

Zur Betonung bemerken wir, dass falls T (durch B) nach oben beschränkt ist, so ist $-T$ (durch $-B$) nach unten beschränkt. Insbesondere sind “untere Schranken” im Sinne der Anordnung nicht notwendig “klein” im Sinne des Absolutbetrags:

Definition 1.7. Wir setzen für alle $a \in \mathbb{K}$:²

$$|a| := \max\{a, -a\} = \begin{cases} a & \text{falls } a \geq 0 \\ -a & \text{falls } a < 0 \end{cases} \quad (1.11)$$

²Im Allgemeinen bezeichnet für eine endliche nicht-leere Teilmenge X eines angeordneten Körpers $\max(X)$ das grösste Element von X .

§1. DIE REELLEN ZAHLEN

Dann gelten die Regeln:

Lemma 1.8.

$$\begin{aligned}
 &\text{Multiplikativitat: } \forall a, b \in \mathbb{K} : |a \cdot b| = |a| \cdot |b|, \\
 &\text{Dreiecksungleichung: } \forall a, b \in \mathbb{K} : |a + b| \leq |a| + |b|, \\
 &\text{“Umgekehrte” Dreiecksungl.: } \forall a, b \in \mathbb{K} : |a - b| \geq ||a| - |b||, \\
 &\text{ausserdem: } |a| = 0 \Leftrightarrow a = 0
 \end{aligned} \tag{1.12}$$

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 Abschatzung zweier Zahlen auch eine Abschatzung ihres Maximums liefert:

$$\left. \begin{array}{l} a + b \leq |a| + |b| \\ -(a + b) \leq |a| + |b| \end{array} \right\} \Rightarrow \max\{a + b, -(a + b)\} \leq |a| + |b| \tag{1.13}$$

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

- Fur die Beweise der Analysis, sowie ihre Anwendungen in der Physik sind Abschatzungen uber den Absolutbetrag von zentraler Bedeutung. Den feinen Unterschied zum Groenvergleich direkt uber die Anordnung werden wir noch zu schatzen lernen.
- Eine fundamental angeordnete physikalische Megroe ist die Zeit.
- Ebenso wichtig ist die Gesamtenergie eines physikalischen Systems im Sinne der Anordnung nach unten beschrankt; sonst droht (bei Kopplung an externe Freiheitsgrade) Instabilitat.
- Eine Grundidee der modernen Physik ist, dass letztlich alle physikalischen Groen absoluten Schranken unterliegen. (Beispiel: Lichtgeschwindigkeit, Plancksches Wirkungsquantum). In der klassischen und mathematischen Idealisierung hingegen spielt fur die Entwicklung der Unendlichkeit der folgende Begriff eine wichtige Rolle:

Definition 1.9. Ein angeordneter Korper $\mathbb{K}$ heist *Archimedisch* falls $\mathbb{N} \subset \mathbb{K}$ nicht nach oben beschrankt ist. M.a.W.

$$\forall x \in \mathbb{K} : \exists N \in \mathbb{N} \text{ s.d. } N > x \tag{1.14}$$

Bemerkungen. · Der Korper der rationalen Zahlen $\mathbb{Q}$ ist Archimedisch. (Fur $x = \frac{p}{q}$ $p > 0$ gilt mit $N = p + 1$: $N > x$.)

- Die Konstruktion von angeordneten nicht-Archimedischen Korpern leuchtet keinem Anfanger sofort ein, weshalb die Wichtigkeit dieses Archimedischen Axioms von allen hier gegebenen Definitionen am schwierigsten einzusehen ist.
- Physikalisch gesehen ist die Unbeschranktheit der naturlichen Zahlen aber sehr naturlich (fur jeden einzelnen Physiker etwa symbolisiert durch den Herzschlag, fur die Menschheit durch die Revolution der Erde um die Sonne).

- Die Umformulierung zu der Aussage: Für alle $B, \epsilon \in \mathbb{K}$, $\epsilon > 0$ 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.
- Das Archimedische Axiom dient dazu, dass wir unsere Zahlenbereiche nicht zu weit erweitern, indem es Anordnungsunendlichkeiten, welche grösser sind als die natürlichen Zahlen, ausschliesst. (Wie wir wissen und evtl. sehen werden, sind allerdings im Sinne der Mächtigkeit von Mengen die reellen Zahlen eine grössere Unendlichkeit als die natürlichen...)

Proposition 1.10. *Sei $\mathbb{K}$ ein Archimedisch angeordneter Körper, und $q, Q \in \mathbb{K}$.*

- (i) *Ist $Q > 1$ so gibt es zu jedem $B \in \mathbb{K}$ 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 1.5 liefert $Q^n \geq 1 + nx$. Weil $\mathbb{K}$ 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 $B = \epsilon^{-1}$. □

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 (1.3). Es ist leicht einzusehen, dass die Lösbarkeit dieser Gleichungen äquivalent ist zur Aussage:

$$\forall(a', b') \in \mathbb{K} \times \mathbb{K}^* : \exists x : b' \cdot x + a' = 0 \quad (1.15)$$

In der Linearen Algebra lernt man Vektorräume über Körpern als die natürlichen Kategorien kennen, in denen sich eine sinnvolle Theorie der linearen Gleichungen als Verallgemeinerung von (1.15) und ihrer Lösungen entwickeln lässt. (Die Vektorraumstruktur bei physikalischen Messgrössen passt damit in wunderbarer Weise zur Idee, dass Lineare Algebra als mathematische Begleitvorlesung zur Theoretischen Physik I ausreicht!)

· Andererseits haben, wie hinreichend bekannt, gewisse natürlich auftretende nicht-lineare Aufgaben wie z.B.

$$x^2 = n, \quad n \in \mathbb{N} \setminus \{m^2 \mid m \in \mathbb{N}\} \quad (1.16)$$

keine Lösung in $\mathbb{Q}$.³

· Eine Möglichkeit zur Abhilfe ist die weitere Erweiterung der rationalen Zahlen um die fehlenden Lösungen von Gleichungen wie (1.16) und deren Verwandten höheren Grades. 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 (1.16) *näherungsweise* zu lösen (auch wenn

³Zur Erinnerung: Jede rationale Zahl lässt sich in eindeutiger Weise als $x = \frac{p}{q}$ darstellen mit $(p, q) \in \mathbb{Z} \times \mathbb{N}$ 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 $n = p^2 \in \{m^2 \mid m \in \mathbb{N}\}$.

§1. DIE REELLEN ZAHLEN

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 (3.24) mit dem Beweis der Konvergenz des Näherungsverfahrens wieder schliessen werden.

· Es sei $x_0 \in \mathbb{Q}$ eine “näherungsweise” Lösung der Gleichung $x^2 = 2$, d.h. $\delta_0 := x_0^2 - 2 \in \mathbb{Q}$ 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

$$x_1 := x_0 - \frac{\delta_0}{2x_0} \quad (1.17)$$

so gilt

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

$$\delta_1 := x_1^2 - 2 = \left(x_0 - \frac{\delta_0}{2x_0}\right)^2 - 2 = x_0^2 - 2 - \delta_0 + \left(\frac{\delta_0}{2x_0}\right)^2 = \frac{\delta_0}{x_0^2} \cdot \frac{\delta_0}{4} \quad (1.18)$$

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 auf verschiedene äquivalente Weisen erfassen. Manche davon resonieren besonders gut mit der physikalischen Idee der Approximation (mittels Absolutbetrag), und wir werden sie im nächsten Kapitel 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.

Definition 1.11. Sei $\mathbb{K}$ ein angeordneter Körper, und $T \subset \mathbb{K}$ eine Teilmenge. Eine Zahl $B \in \mathbb{K}$ heisst *kleinste obere Schranke* von T falls

- (i) B eine obere Schranke ist und
(ii) Für jede andere obere Schranke B' von T gilt, dass $B' \geq B$.

Ein angeordneter Körper $\mathbb{K}$ heisst (*anordnungs-*)*vollständig*, falls jede nicht-leere nach oben beschränkte Menge eine kleinste obere Schranke besitzt.

(Die analoge Definition über grösste untere Schranken ist 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 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.
- Man kann auch zeigen, dass eine etwaige kleinste obere Schranke dieser Menge eine Lösung der Gleichung $x^2 = 2$ wäre, welche bekanntlich nicht rational ist.

Theorem 1.12. (i) *Es existiert ein bis auf Isomorphismus eindeutig bestimmter vollständiger angeordneter Körper, genannt der Körper der reellen Zahlen $\mathbb{R}$.*

(ii) *$\mathbb{R}$ ist Archimedisch.*

(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 über diejenigen *rationalen* Zahlen, welche strikt grösser oder kleiner sind.

· Formal ist ein *Dedekindscher Schnitt* 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. $\forall y < x : y \in A \Rightarrow x \in A$ und

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

Beispiele für solche Schnitte sind für $r \in \mathbb{Q}$: $A_r = \{x \mid x < r\}$. 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 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) Beachte, dass man für die Definition von “archimedisch” ($\forall x \exists N > x$) braucht, dass jeder angeordnete Körper die ganzen Zahlen enthält.

Angenommen, es gäbe ein $x \in \mathbb{K}$ so, dass $x \geq N \forall N \in \mathbb{N}$. Dann wäre x eine obere Schranke für $\mathbb{N}$, und es gäbe eine kleinste obere Schranke, x_0 . $\Rightarrow x_0 > N \forall N \in \mathbb{N}$. $\Rightarrow x_0 - 1 > N \forall N \in \mathbb{N}$. Dann aber wäre $x_0 - 1 < x_0$ ebenfalls eine obere Schranke, kleiner als x_0 .

(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

§1. DIE REELLEN ZAHLEN

genannten Eigenschaften *voraus*, und formulieren die Vollständigkeit weiter unten auf so viele verschiedenen Weisen um, wie es für unsere Zwecke nützlich ist.

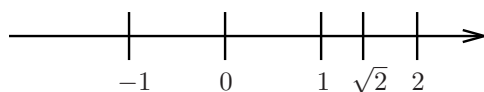
- Es ist letztlich dem Genie Isaac Newtons die Idee zu verdanken, dass die idealisierte Einführung der reellen Zahlen es nicht nur erlaubt, (gewisse, aber bei weitem (noch) nicht alle, siehe nächster Abschnitt) algebraische Gleichungen zu lösen, sondern auch, Differentialgleichungen zu formulieren, deren Lösungen (“Integrale”) die Bewegungen von Massenpunkten, die Ausbreitung von Wellen und klassischen Feldern, und dergleichen mehr, in einer geeigneten Näherung hervorragend beschreiben.
- Das Ziel dieser Vorlesung ist der theoretische Nachweis der Existenz solcher Lösungen in diesem Rahmen sowie die Entwicklung geeigneter Lösungsmethoden.

Definition 1.13. · Für eine nach oben beschränkte nichtleere Teilmenge $T \subset \mathbb{R}$ nennen wir die kleinste obere Schranke von T das Supremum von T ($\sup T$). Falls $\sup T \in T$, so ist $\sup T = \max T$, dem Maximum von T .

· Für eine nach unten beschränkte nichtleere Teilmenge $T \subset \mathbb{R}$ nennen wir die größte untere Schranke das Infimum von T ($\inf T$). Falls $\inf T \in T$, so ist $\inf T = \min T$, dem Minimum von T .

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.



Definition 1.14. Besonders interessant und als Definitionsbereiche von Funktionen wichtige Teilmengen von $\mathbb{R}$ sind die Intervalle. Für $a < b$ unterscheiden wir:

abgeschlossen: $[a, b] := \{x | x \geq a, b \geq x\}$

“halboffen”: $\begin{cases} [a, b) := \{x | x \geq a, b > x\} \\ (a, b] := \{x | x > a, b \geq x\} \end{cases}$

offen: $(a, b) := \{x | x > a, b > x\}$

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.⁴ Beschränkte und abgeschlossene Intervalle heissen *kompakt*.

Proposition 1.15. · 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 $x = a^{1/n}$.

- Aus $b \geq a > 0$ folgt $b^{1/n} \geq a^{1/n}$.
- Es gilt $|x| = (x^2)^{1/2} \forall x > 0$.

⁴Man beachte auch, dass dem Wort “abgeschlossen” hier ein recht anderer Sinn gegeben wird als in (β). Die ganze Menge $\mathbb{R} = (-\infty, \infty)$ ist sowohl offen als auch abgeschlossen. Die beschränkten halboffenen Intervalle sind weder offen noch abgeschlossen. Für manche Zwecke ist es nützlich, auch ein-Punkt Mengen als Intervalle $\{a\} = [a, a]$ zuzulassen.

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

§ 2 Die komplexen Zahlen

Lemma/Definition 2.1. Die Menge aller Paare $(x, y) \in \mathbb{R} \times \mathbb{R}$ reeller Zahlen, versehen mit den Operationen:

$$\text{Addition: } (x_1, y_1) + (x_2, y_2) := (x_1 + x_2, y_1 + y_2)$$

$$\text{Multiplication: } (x_1, y_1) \cdot (x_2, y_2) := (x_1x_2 - y_1y_2, x_1y_2 + x_2y_1)$$

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

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

$$(x, y) \cdot \left(\frac{x}{x^2 + y^2}, \frac{-y}{x^2 + y^2} \right) = (1, 0) = 1_{\mathbb{C}} \quad (2.1)$$

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. für $\mathbb{K}$ einen angeordneten Körper $x^2 \geq 0 \forall x \in \mathbb{K}$, und daraus folgt $x^2 + 1 \geq 1$, mithin hat die Gleichung $z^2 + 1 = 0$ keine Lösung in $\mathbb{K}$.

· In $\mathbb{C}$ hingegen gilt

$$(0, 1)^2 = (-1, 0) = -1_{\mathbb{C}} \quad (2.2)$$

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

$$z^n + a_{n-1}z^{n-1} + \cdots + a_1z + a_0 = 0 \quad a_i \in \mathbb{C}, n \in \mathbb{N} \quad (2.3)$$

eine Lösung $z \in \mathbb{C}$ (Mit Multiplizität gezählt gibt es dann n Lösungen.) (Fundamentalsatz der Algebra.) 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 Sinner der Anordnung auf $\mathbb{R}$.”

· Definieren wir nun $i := (0, 1) \in \mathbb{C}$ und denken uns $\mathbb{R}$ als in $\mathbb{C}$ eingebettet via

$$\mathbb{R} \ni x \mapsto (x, 0) \in \mathbb{C} \quad (2.4)$$

(d.h., wir identifizieren $x \in \mathbb{R}$ mit $(x, 0) \in \mathbb{C}$), dann können wir jede komplexe Zahl in eindeutiger Weise schreiben als

$$z = (x, y) = x + iy \quad x, y \in \mathbb{R} \subset \mathbb{C} \quad (2.5)$$

§2. DIE KOMPLEXEN ZAHLEN

und “wie gewöhnlich” rechnen.

· M.a.W. ist $\mathbb{C}$ ein 2-dimensionaler Vektorraum über $\mathbb{R}$ mit Basis $1_{\mathbb{C}} = (1, 0)$ und $i = (0, 1)$. Allerdings wollen wir uns möglichst bald daran gewöhnen, komplexe Zahlen als eine Einheit aufzufassen.

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

$$x + iy = z \mapsto \bar{z} = x - iy \quad (2.6)$$

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

$$\operatorname{Re}(z) := \frac{z + \bar{z}}{2} = x, \quad \operatorname{Im}(z) := \frac{z - \bar{z}}{2i} = y \quad (2.7)$$

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

$$|z| := \begin{cases} (\bar{z}z)^{1/2} & z \neq 0 \\ 0 & z = 0 \end{cases} \quad (2.8)$$

welcher wegen Prop. 1.15 auf $\mathbb{R}$ mit (1.11) übereinstimmt. Es gelten $|z| = 0 \Leftrightarrow z = 0$, Abschätzungen der Art

$$|\operatorname{Re}(z)| \leq |z| \quad |\operatorname{Im}(z)| \leq |z| \quad (2.9)$$

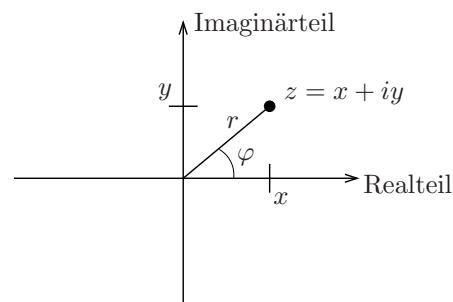
sowie insbesondere die Dreiecksungleichung

$$|z + w| \leq |z| + |w| \quad (2.10)$$

(vgl. Lemma 1.8): Aus den Rechenregeln folgt $|z + w|^2 = |z|^2 + |w|^2 + 2\operatorname{Re}(\bar{z}w) \leq |z|^2 + |w|^2 + 2|z| \cdot |w| = (|z| + |w|)^2$ unter Benutzung von (2.9). Da für $a \geq 0$, $b \geq 0$ aus $a^2 \geq b^2$ folgt, dass auch $a \geq b$ (vgl. 1.15), impliziert dies (2.10).

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 $z \neq 0$

$$z = r \left(\underbrace{\frac{x}{r}}_{=:c} + i \underbrace{\frac{y}{r}}_{=:s} \right) \quad (2.11)$$

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

$$z = r(\cos \varphi + i \sin \varphi) \quad (2.12)$$

Mithilfe der Additionstheoreme für $\sin$ und $\cos$ folgt 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

$$z \cdot w = r \cdot s(\cos(\varphi + \chi) + i \sin(\varphi + \chi)) \quad (2.13)$$

Achtung: (2.13) ist i.A. *nicht* die Darstellung von $z \cdot w$ in der Form (2.12), denn $\varphi + \chi$ kann ausserhalb des Intervalls $[0, 2\pi)$ liegen. Strenggenommen machen diese Betrachtungen auch erst nach der Definition der Winkelfunktionen und dem Nachweis ihrer Periodizität Sinn.

· Daraus bzw. durch vollständige Induktion folgt dann z.B. auch die Formel von deMoivre

$$z^n = r^n(\cos(n\varphi) + i \sin(n\varphi)) \quad (2.14)$$

sowie die $\mathbb{C}$ -Version von Prop. 1.15, einem Spezialfall von (2.3).

Proposition 2.2. *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

$$z_k = r^{1/n} \left(\cos \left(\frac{\varphi}{n} + \frac{2\pi k}{n} \right) + i \sin \left(\frac{\varphi}{n} + \frac{2\pi k}{n} \right) \right) \quad (2.15)$$

für alle $k = 0, 1, \dots, n-1$ die gegebene Gleichung. Dass dies alle Lösungen sind, folgt aus der aus der Algebra 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 folgenden § 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 2.3. Eine Teilmenge $T \subset \mathbb{C}$ heisst beschränkt, falls eine Schranke $B \in \mathbb{R}$ existiert mit

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

§ 2. DIE KOMPLEXEN ZAHLEN

· Zuletzt erwähnen wir noch als Analogon der Intervalle als “elementare” Definitionsbereiche von Funktionen im Komplexen die Kreisscheiben vom Radius $R > 0$ um den Punkt $z \in \mathbb{C}$:

$$D_R(z) := \{w \in \mathbb{C} \mid |z - w| < R\} \quad (2.16)$$

↑
Beachte hier $\neq$. Dies entspricht
den offenen Intervallen.

KAPITEL 2

KONVERGENZ

In der Analysis formuliert und löst man mathematische Probleme (die wir uns als aus der Physik stammend denken) unter Ausnutzung der Vollständigkeit im Rechenbereich der reellen oder komplexen Zahlen anstatt mit expliziter Algebra oder durch Vorzeigen. Dazu ist allerdings, wie bereits angedeutet, unsere bisherige Formulierung via Supremumseigenschaft zu glatt (unter anderem macht sie ja im Komplexen keinen Sinn). In diesem Kapitel entwickeln wir zunächst aus der Supremumseigenschaft die Grundlagen der Grenzprozesse, die sich anschliessend als wesentlich allgemeiner und bequemer herausstellen. Die Umformulierung benutzt die typischen “zu jedem $\epsilon > 0$ ”-enthaltenden Wendungen, deren Banalität wir an folgendem häufig eingesetzten Argument illustrieren wollen.

Lemma 3.-1. *Über eine reelle Zahl $x \in \mathbb{R}$ sei bekannt, dass für jedes $\epsilon > 0$, $x \in [-\epsilon, \epsilon]$ (m.a.W., für alle $\epsilon > 0$ ist $|x| \leq \epsilon$). Dann ist $x = 0$.*

Beweis. Wir zeigen zunächst, dass x nicht positiv ist: Wäre $x > 0$, so gälte für $\epsilon_* := \frac{x}{2}$: $\epsilon_* > 0$ und $x > \epsilon_*$, d.h. $x \notin [-\epsilon_*, \epsilon_*]$, im Widerspruch zur Voraussetzung des Lemmas.

· Wäre hingegen $x < 0$, so gälte für $\epsilon_* := \frac{-x}{2}$ wieder $\epsilon_* > 0$ und $x < -\epsilon_*$, d.h. $x \notin [-\epsilon_*, \epsilon_*]$.

· Es bleibt damit nur $x = 0$. □

· Tatsächlich genügt es wegen der Archimedischen Eigenschaft von $\mathbb{R}$, wenn man in solchen Aussagen das beliebige $\epsilon > 0$ durch “ $\frac{1}{N}$ für jedes $N \in \mathbb{N}$ ” ersetzt.

· Man überlege sich auch an dieser Stelle, dass man in dieser Aussage das abgeschlossene Intervall auch durch das offene ersetzen kann.

§ 3 Folgen reeller Zahlen

Definition 3.1. Sei $X \neq \emptyset$ eine Menge. Unter einer Folge in X (oder auch mit Werten in X oder auch X -ige Folge) versteht man eine Abbildung

$$a : \mathbb{N} \rightarrow X \quad n \mapsto a(n) = a_n \in X \quad (3.1)$$

von der Menge der natürlichen Zahlen in X . Man schreibt auch $(a_n)_{n \in \mathbb{N}}$, $(a_n)_{n=1,2,\dots}$, $(a_n) \subset X$.

· 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 die Menge $\{a_n \mid n \in \mathbb{N}\} \subset X$ aufzufassen. Andere Aussagen betreffen die Abbildung (3.1).

Definition 3.2. (i) Eine Folge (x_n) reeller Zahlen heisst beschränkt, wenn eine Schranke $B \geq 0$ existiert so, dass $|x_n| \leq B$ für alle $n \in \mathbb{N}$.
 (ii) (x_n) heisst monoton wachsend, falls für alle $n, m \in \mathbb{N}$

$$n > m \Rightarrow x_n \geq x_m$$

(Offenbar ist dies gleichwertig zu $x_{n+1} \geq x_n \forall n \in \mathbb{N}$.)

(iii) (x_n) heisst streng monoton wachsend falls auf der rechten Seite $x_n \geq x_m$ (d.h. $x_n > x_m$) steht. Sie heisst (streng) monoton fallend, falls $x_n \leq x_m$ (bzw. $x_n < x_m$).

Definition 3.3. 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 Zahlen via

$$b_k := a_{n_k} \quad (3.2)$$

gewonnen worden ist. (Das heisst, man zählt nur einen Teil der Folgenglieder von (a_n) ab, aber unter Berücksichtigung von deren ursprünglichen Reihenfolge.)

Die Idee ist nun, reelle Zahlen nicht explizit (wie etwa durch Dedekindsche Schnitte) anzugeben, sondern ihnen durch *sukzessive Approximation*, nämlich durch Folgen, allmählich immer näher zu kommen. Dies ist voll an der physikalischen Intuition ausgerichtet.

Definition 3.4. 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, sodaß

$$\begin{aligned} & |x_n - x| < \epsilon \quad \forall n \geq N_\epsilon \\ \Leftrightarrow & x - x_n \in (-\epsilon, +\epsilon) \quad \forall n \geq N_\epsilon \\ \Leftrightarrow & x_n \in (x - \epsilon, x + \epsilon) \quad \forall n \geq N_\epsilon \end{aligned} \quad (3.3)$$

(Wir sagen auch: (x_n) konvergiert “gegen”, oder “mit Grenzwert” x .) Eine nicht-konvergente reelle Folge heisst “divergent”.

Bemerkungen. · Es ist oft nützlich, leichte Variationen der Definition als gleichwertig zu erkennen, wie z.B.: “falls ein $x \in \mathbb{R}$ und ein $c > 0$ existieren so, dass $\forall \epsilon > 0$ ein $N_\epsilon \in \mathbb{N}$ existiert so, dass

$$|x - x_n| \leq c \cdot \epsilon \quad \forall n > N_\epsilon” \quad (3.4)$$

· Falls (x_n) konvergiert, so ist der Grenzwert eindeutig:

Beh.: Falls (x_n) gegen x und y konvergiert, dann gilt $x = y$.

Bew.: Für $\epsilon > 0$ seien N_ϵ und M_ϵ so, dass

$$\begin{aligned} & |x_n - x| < \epsilon \quad \forall n \geq N_\epsilon \\ \text{und} & |x_n - y| < \epsilon \quad \forall n \geq M_\epsilon \end{aligned} \quad (3.5)$$

Dann gilt mit einem $n_0 \geq \max\{N_\epsilon, M_\epsilon\}$:

$$|x - y| \leq |x - x_{n_0}| + |x_{n_0} - y| < 2\epsilon \quad (3.6)$$

§3. FOLGEN REELLER ZAHLEN

Wegen Lemma 3.1 folgt aus diesem Schluss $x = y$. □

· Für eine konvergente Folge macht also die Schreibweise

$$\lim_{n \rightarrow \infty} x_n = x \text{ oder auch } x_n \xrightarrow[n \rightarrow \infty]{} x \quad (3.7)$$

Sinn.

Definition 3.5. Eine Nullfolge ist eine konvergente Folge (*p.t.* in $\mathbb{R}$) mit Grenzwert 0.

Lemma 3.6. 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.

Beispiel 3.7. Wichtige Folgengrenzwerte:

- (i) Für jedes positive $s \in \mathbb{Q}$ ist $\lim_{n \rightarrow \infty} \frac{1}{n^s} = 0$
- (ii) Für $a > 0$ gilt $\lim_{n \rightarrow \infty} a^{1/n} = 1$
- (iii) $\lim_{n \rightarrow \infty} n^{1/n} = 1$
- (iv) Für $0 < q < 1$ ist $\lim_{n \rightarrow \infty} q^n = 0$
- (v) Für alle $k \in \mathbb{N}$ und $x > 1$ ist $\lim_{n \rightarrow \infty} \frac{n^k}{x^n} = 0$

Bew.: (i) Für jedes $\epsilon > 0$ gilt mit $N_\epsilon > \frac{1}{\epsilon^{1/s}}$ (ein solches existiert nach dem Archimedischen Axiom) für alle $n \geq N_\epsilon$:

$$\frac{1}{n^s} \leq \frac{1}{N_\epsilon^s} < \left(\frac{1}{\epsilon^{1/s}} \right)^{-s} = \epsilon \quad (3.8)$$

(ii) Wir behandeln zunächst den Fall $a > 1$: Nach der Bernoullischen Ungleichung 1.5 gilt

$$\begin{aligned} a &= \left(1 + \underbrace{a^{1/n} - 1}_{>0} \right)^n \geq 1 + n(a^{1/n} - 1) \\ \Rightarrow \quad 0 &< a^{1/n} - 1 \leq \frac{a - 1}{n} \end{aligned} \quad (3.9)$$

Die Aussage folgt dann aus (i) und dem sogenannten Sandwich-Theorem:

Lemma 3.8. Es sei (x_n) eine reelle Folge und (A_n) und (B_n) konvergente reelle Folgen mit

$$\lim A_n = \lim B_n =: x \quad \text{und} \quad A_n \leq x_n \leq B_n \quad \forall n \quad (3.10)$$

Dann gilt $\lim_{n \rightarrow \infty} x_n = x$.

Beweis. Übungsaufgabe. □

Bemerkungen. · Es genügt 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 $n \geq N$.

· Es sei auch noch einmal darauf hingewiesen, dass für die Abschätzungen im Gross-
teil der Aussagen und Beweise die Unterscheidung zwischen $<$ und $\leq$ unerheblich
ist, s. auch (3.4).

· 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 folgt
 $\lim A_n \leq \lim B_n$, aber im Allgemeinen gilt nicht $<$! (Konkret: $a_n = 0 \forall n, b_n = 1/n$.)

Bevor wir mit den Beispielen fortfahren, halten wir die folgenden Rechenregeln
fest:

Proposition 3.9. Für reelle Folgen (x_n) und (y_n) gelte $x_n \rightarrow x$ und $y_n \rightarrow y$. Dann
gilt

$$(\alpha) \lim(x_n + y_n) = x + y$$

$$(\beta) \lim x_n \cdot y_n = x \cdot y$$

$$(\gamma) \lim |x_n| = |x|$$

$$(\delta) \text{ falls } y \neq 0, \text{ so sind fast alle } y_n \neq 0 \text{ und } \lim \frac{x_n}{y_n} = \frac{x}{y}$$

Beweis. siehe Lehrbücher. □

Bemerkungen. Natürlich gelten hier keine Umkehrungen. Insbesondere folgt aus
 $|x_n| \rightarrow |x|$ nicht $x_n \rightarrow x$, siehe etwa 3.11.

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 Nachweis von (iii) betrachten wir zunächst, wieder mit der Bernoullischen
Ungleichung

$$\begin{aligned} n^{1/2} &= (1 + n^{1/2n} - 1)^n \geq 1 + n(n^{1/2n} - 1) \\ \Rightarrow 0 &\leq n^{1/2n} - 1 \leq \frac{n^{1/2} - 1}{n} = \frac{1}{n^{1/2}} - \frac{1}{n} \xrightarrow{n \rightarrow \infty} 0 \end{aligned} \quad (3.11)$$

Also gilt $\lim_{n \rightarrow \infty} n^{1/2n} = 1$, und wegen Regel (β) folgt

$$\lim_{n \rightarrow \infty} n^{1/n} = \lim_{n \rightarrow \infty} (n^{1/2n} \cdot n^{1/2n}) = 1 \quad (3.12)$$

· Beispiel (iv) ist genau die Aussage von 1.10.

· Zum Beweis von (v) betrachten wir, schon wieder mit Bernoulli:

$$\begin{aligned} x^n &\geq 1 + n(x - 1) > n^{1/2} \cdot n^{1/2} \underbrace{(x - 1)}_{>0} \\ \Rightarrow \frac{n^{1/2}}{x^n} &< \frac{1}{n^{1/2}(x - 1)} \xrightarrow{n \rightarrow \infty} 0 \end{aligned} \quad (3.13)$$

Dann folgt die Aussage wieder mit (β) :

$$\frac{n^k}{x^n} = \left(\frac{n^{1/2}}{\underbrace{(x^{1/2k})^n}_{>1}} \right)^{2k} \xrightarrow{n \rightarrow \infty} 0 \quad (3.14)$$

§3. FOLGEN REELLER ZAHLEN

- Alle diese Folgen (i)–(v), sowie die Rechenregeln, sind zwar sehr nützlich, illustrieren aber gerade *nicht* die wesentliche Idee, Folgen zum Herzeigen “neuer” reeller Zahlen (insbesondere solche in $\mathbb{R} \setminus \mathbb{Q}$) zu benützen. Dafür müssen wir aus der Vollständigkeit 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

$$x_{n+1} = x_n - \frac{x_n^2 - 2}{2x_n} = \frac{x_n}{2} + \frac{1}{x_n} \quad (3.15)$$

(vgl. (1.17)), welche für beliebigen Startwert $x_0 > \sqrt{2}$ die Eigenschaft hat, dass

$$\lim_{n \rightarrow \infty} x_n = \sqrt{2}, \quad (3.16)$$

was wir als Konsequenz aus unserem ersten grenzwertfreien Konvergenzkriterium bald beweisen. Zunächst noch eine kleine Übung:

Lemma 3.10. *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*

$$\lim_{k \rightarrow \infty} x_{n_k} = \lim_{n \rightarrow \infty} x_n \quad (3.17)$$

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 $|x_{n_k} - x| \leq \epsilon \ \forall k \geq N_\epsilon$.

· 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 der Teilfolgen. $\square$

Beispiel 3.11. Es genügt natürlich nicht, wenn irgendeine Teilfolge konvergiert: Die Folge $(a_n)_{n \in \mathbb{N}} = (-1, 1, -1, \dots)$ mit

$$a_n := (-1)^n = \begin{cases} -1 & n \text{ ungerade} \\ +1 & n \text{ gerade} \end{cases} \quad (3.18)$$

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 sind alle Teilfolgen $(a_{n_k})_{k \in \mathbb{N}}$ für $n_k = p^k$ Nullfolgen (die meisten anderen sowieso), aber $(a_p)_{p \text{ prim}}$ ist konstant 1 und (a_n) selbst ist divergent.

· Es ist instruktiv sich zu überlegen, dass das Weglassen endlich vieler Folgenglieder das Konvergenzverhalten nicht ändert. Als Beispiel hierfür betrachten wir für $x \in \mathbb{R}$ die Folge

$$a_n = \frac{x^n}{n!} \quad (3.19)$$

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 $n \geq N_0$:

$$\begin{aligned} |a_n| &= \left| a_{N_0} \cdot \frac{x^{n-N_0}}{(N_0+1) \cdots (n-1) \cdot n} \right| \\ &= |a_{N_0}| \cdot \frac{|x|}{N_0+1} \cdots \frac{|x|}{n-1} \cdot \frac{|x|}{n} \\ &\leq |a_{N_0}| \cdot 2^{N_0-n} \rightarrow 0 \quad \text{für } n \rightarrow \infty \end{aligned} \quad (3.20)$$

· Mit ähnlichen Überlegungen zeigt man, dass für $x > 1$ die Folge

$$a_n = \frac{x^{n^2}}{n!} \quad (3.21)$$

divergiert (und zwar bestimmt gegen ∞). (Übungsaufgabe)

Wir wollen nun daran gehen, die Supremumseigenschaft von $\mathbb{R}$ in die Folgen-sprache zu übersetzen. (Für den Monotonie-Begriff siehe Def. 3.2.)

Theorem 3.12. *Jede monoton wachsende und nach oben beschränkte Folge reeller Zahlen ist konvergent.*

Bemerkungen. Ditto jede monoton fallende und nach unten beschränkte Folge. Wenn wir “bestimmte Divergenz” mit Grenzwert $\pm\infty$ zulassen, so kann die Beschränktheit fallengelassen werden. Unter Voraussetzung des Archimedischen Axioms ist das Prinzip der monotonen Konvergenz äquivalent zur Anordnungs-Vollständigkeit (Übungsaufgabe).

Beweis von 3.12. Sei (x_n) eine monoton wachsende nach oben beschränkte reelle Folge. Die Menge $T = \{x_n | n \in \mathbb{N}\} \subset \mathbb{R}$ ist dann nichtleer und nach oben beschränkt. Nach der Supremumseigenschaft (Anordnungsvollständigkeit) existiert eine kleinste obere Schranke,

$$s := \sup\{x_n | n \in \mathbb{N}\} \quad (3.22)$$

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 $n \geq N_\epsilon$

$$\begin{array}{c} s \geq x_n \geq x_{N_\epsilon} > s - \epsilon \\ \uparrow \\ s \text{ ist obere Schranke} \end{array} \quad (3.23)$$

d.h. $|x_n - s| < \epsilon, \forall n \geq N_\epsilon$. Dies impliziert die Behauptung. $\square$

§3. FOLGEN REELLER ZAHLEN

Bew.: [von (3.16)] Aus den Betrachtungen um (1.18) folgt, dass (x_n) monoton fällt und nach unten beschränkt ist. Daher existiert jedenfalls $x = \lim x_n$, 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. 3.10).
- Andererseits gilt wegen $(\alpha), (\delta)$:

$$\begin{aligned} x = \lim y_n &= \lim \frac{x_n}{2} + \lim \frac{1}{x_n} = \frac{x}{2} + \frac{1}{x} \\ \Rightarrow x &= \sqrt{2} \end{aligned} \quad (3.24)$$

Die zweite wichtige Methode zur Vorstellung reeller Zahlen ist:

Theorem/Definition 3.13. *Eine Intervallschachtelung ist eine Folge von abgeschlossenen und beschränkten (kurz: kompakten) Intervallen $([a_n, b_n])_{n \in \mathbb{N}}$ mit*

$$\begin{aligned} [a_{n+1}, b_{n+1}] &\subset [a_n, b_n] \quad \forall n \\ \text{und } \lim_{n \rightarrow \infty} (b_n - a_n) &= \lim |b_n - a_n| = 0 \end{aligned} \quad (3.25)$$

(i) Zu jeder Intervallschachtelung $([a_n, b_n])$ existiert genau eine reelle Zahl x , die in allen $[a_n, b_n]$ enthalten ist, d.h.

$$\bigcap_{n=1}^{\infty} [a_n, b_n] = \{x\} \quad (3.26)$$

(ii) Zu jeder reellen Zahl x existieren (sehr viele!) Intervallschachtelungen mit Durchschnitt x .

Beweis. (i) Die Folge (a_n) ist monoton wachsend und nach oben beschränkt (durch irgendein b_n), daher konvergent wegen 3.12. Setze $x = \lim a_n$. Dann gilt $x \geq a_n \forall n$. Für $\epsilon > 0$ sei nun N_ϵ derart, dass gleichzeitig

$$\left. \begin{aligned} b_n - a_n &< \epsilon \\ x - a_n &< \epsilon \end{aligned} \right\} \forall n \geq N_\epsilon \quad (3.27)$$

(Ersteres lässt sich erreichen wegen $\lim b_n - a_n \rightarrow 0$, letzteres wegen $\lim a_n = x$.) Dann gilt $\forall n \geq N_\epsilon$

$$-\epsilon < a_n - x \leq b_n - x \leq b_n - a_n < \epsilon \quad (3.28)$$

Es folgt $\lim b_n = x$, und wegen der Monotonie von (b_n) , dass $x \leq b_n$, also $x \in [a_n, b_n] \forall n$. Eindeutigkeit von x , und (ii), als Übung. $\square$

Warnung: Das Prinzip gilt *nicht* mit offenen Intervallen. Z.B. ist $\bigcap_{n=1}^{\infty} (0, \frac{1}{n}) = \emptyset$. Besonders mächtig ist

Theorem 3.14. *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

$$\{x_n | n \in \mathbb{N}\} \subset [-B, B] \quad (3.29)$$

· Zur Identifikation eines mutmasslichen Grenzwertes definieren wir rekursiv eine Intervallschachtelung $([a_k, b_k])_{k \in \mathbb{N}}$ so, dass für jedes k gilt:

$$\text{Bed}(k) : [a_k, b_k] \text{ enthält unendlich viele } x_n \quad (3.30)$$

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.

· $[a_{k+1}, b_{k+1}] \subset [a_k, b_k]$ ist klar 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 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 Punkt 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.: $x_{n_k} \xrightarrow{k \rightarrow \infty} x$

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 $|x - x_{n_k}| < \epsilon \forall k \geq K$. $\square$

Bemerkungen. Da wir im Beweis stets das obere Halbintervall bevorzugt haben, erfüllt das resultierende x :

$$\begin{aligned} x &= \sup\{y \mid \text{Es existiert eine gegen } y \text{ konvergente Teilfolge von } (x_n)\} \\ &= \max\{y \mid \text{Es existiert eine gegen } y \text{ konvergente Teilfolge von } (x_n)\} \end{aligned} \quad (3.31)$$

Man nennt die Elemente der Menge auf der rechten Seite auch die Häufungswerte der beschränkten Folge (x_n) und deren Supremum x den “limes superior”, geschrieben $x = \limsup(x_n)$. Dazu äquivalente Charakterisierungen sind:

$$\begin{aligned} \limsup(x_n) &= \min\{y \mid \forall \epsilon > 0 \text{ ist } x_n > y + \epsilon \text{ für höchstens endlich viele } n.\} \\ &= \inf\{y \mid x_n > y \text{ für höchstens endlich viele } n.\} \end{aligned} \quad (3.32)$$

Der “limes inferior” ist analog der kleinste Häufungswert.

· Wie bereits angedeutet ist es bei reellen Zahlen und Folgen manchmal bequem, die Beschränktheitsvoraussetzungen in 1.13 bzw. 3.12 fallenzulassen. Man vereinbart dann etwa, dass für eine nach oben (unten) unbeschränkte nicht-leere Menge T $\sup T = +\infty$ ($\inf T = -\infty$). (Für die leere Menge folgen daraus konsequenterweise $\inf(\emptyset) = +\infty$ und $\sup(\emptyset) = -\infty$.) Damit verallgemeinert man (3.31) zu

§3. FOLGEN REELLER ZAHLEN

Lemma/Definition 3.15. Sei $(x_n) \subset \mathbb{R}$ eine beliebige reelle Folge. Dann ist die Folge $\bar{x}_n := \sup\{x_m \mid m \geq n\}$ monoton fallend, und wir setzen

$$\limsup(x_n) = \lim_{n \rightarrow \infty} \bar{x}_n \in \mathbb{R} \cup \{-\infty, +\infty\} \quad (3.33)$$

Ebenso ist $\underline{x}_n := \inf\{x_m \mid m \geq n\}$ monoton wachsend und wir setzen

$$\liminf(x_n) = \lim_{n \rightarrow \infty} \underline{x}_n \in \mathbb{R} \cup \{-\infty, +\infty\} \quad (3.34)$$

Beweis. Die Monotonie der Folgen $(\bar{x}_n)$ (bzw. $(\underline{x}_n)$) folgt aus der allgemeinen Tatsache, dass für $T \subset \mathbb{R}$ und $S \subset T$ jede obere Schranke von T auch eine obere Schranke von S ist. Daher gilt $\sup S \leq \sup T$ (bzw. analog $\inf S \geq \inf T$). $\square$

Zuletzt erreichen wir das Ziel des am meisten eingesetzten Konvergenzkriteriums.

Definition 3.16. Eine Folge reeller Zahlen (x_n) heisst Cauchy-Folge (CF), falls für jedes $\epsilon > 0$ ein $N_\epsilon \in \mathbb{N}$ existiert derart, dass

$$|x_n - x_m| < \epsilon \quad \text{falls } n \geq N_\epsilon \text{ und } m \geq N_\epsilon \quad (3.35)$$

Theorem 3.17. Eine Folge reeller Zahlen konvergiert dann und nur dann wenn sie eine Cauchy-Folge ist.

Bemerkungen. · Der Vorteil des Cauchy-Kriteriums gegenüber der ursprünglichen Definition 3.4 ist, dass es ohne explizite Erwähnung des Grenzwertes auskommt.

· Das Prinzip der monotonen Konvergenz kommt auch ohne den Grenzwert aus. Allerdings braucht dieses Prinzip die Anordnung von $\mathbb{R}$ in direkter 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 $|x_n - x_{N_\epsilon}| < \epsilon \quad \forall n \geq N_\epsilon$ ausreichen (Übungsaufgabe). Hingegen ist etwa schon die Forderung $|x_{n+1} - x_n| < \epsilon \quad \forall n \geq N_\epsilon$ zu schwach.

Der Beweis des Cauchy-Kriteriums ist mit unseren Vorbereitungen relativ einfach. Wir verteilen ihn auf drei Lemmas.

Lemma 3.18. 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. $|x - x_n| < \frac{\epsilon}{2} \quad \forall n \geq N_\epsilon$. Dann gilt für $n \geq N_\epsilon$ und $m \geq N_\epsilon$,

$$|x_n - x_m| \leq |x_n - x| + |x - x_m| < \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon \quad (3.36)$$

$\square$

Lemma 3.19. 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 $m = N_1$ für $n \geq N_1$: $|x_n| = |x_{N_1} + x_n - x_{N_1}| \leq |x_{N_1}| + 1 \forall n \geq N_1$. Mit $B := \max\{|x_1|, |x_2|, \dots, |x_{N_1-1}|, |x_{N_1}| + 1\}$ gilt dann $|x_n| \leq B \forall n$. $\square$

Lemma 3.20. *Jede Cauchy-Folge ist konvergent.*

Beweis. Nach Lemma 3.19 ist eine Cauchy-Folge (x_n) jedenfalls beschränkt. Nach 3.14 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 $n \geq N_\epsilon$ 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 $n \geq N_\epsilon$:

$$|x_n - x| \leq |x_n - x_{n_{K_\epsilon}}| + |x_{n_{K_\epsilon}} - x| < 2\epsilon \quad (3.37)$$

Unter Beachtung von (3.4) folgt die Behauptung. $\square$

Mit 3.20 und 3.18 ist auch 3.17 bewiesen. $\square$

Abschliessende Bemerkungen; Überabzählbarkeit von $\mathbb{R}$

Es sei noch einmal betont, dass während die Definitionen und Rechenregeln zur Konvergenz auch für Folgen rationaler Zahlen gegolten hätten, die Konvergenzkriterien 3.12 (monotone Konvergenz), 3.17 (Cauchy-Kriterium) sowie die Ergebnisse 3.13 (Intervallschachtelungen) sowie 3.14 (Satz von Bolzano-Weierstrass) wesentlich von der Vollständigkeit der reellen Zahlen abhängen.

· Der Preis dafür, und ein Maß für die “Irrationalität” der reellen Zahlen, ist ihre Überabzählbarkeit.

Erinnerung. Eine Menge X heisst abzählbar, wenn eine injektive Abbildung von X auf $\mathbb{N}$ 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 zu 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 3.21. *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.)

§ 4. METRISCHE RÄUME

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 nicht in (x_n) 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 3.12 existiert $a = \lim_{m \rightarrow \infty} a_m$ und $b = \lim_{m \rightarrow \infty} b_m$, und 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 der beiden Folgen (Übung). 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).

§ 4 Metrische Räume

Im letzten § haben wir das Cauchy-Kriterium für die Konvergenz reeller Folgen aus der Supremumseigenschaft von $\mathbb{R}$ hergeleitet, und dabei angekündigt, dass

(i) für einen angeordneten Körper das Cauchy-Kriterium (zusammen mit dem Archimedischen Axiom) eine zur Supremumseigenschaft äquivalente Charakterisierung der Vollständigkeit der reellen Zahlen liefert, und

(ii) wir das Cauchy-Kriterium nutzen wollen, um den Begriff der Vollständigkeit auf Situationen auszudehnen, in denen eine Anordnung nicht vorhanden oder sogar unmöglich ist.

Die Aussage (i) wollen wir hier nicht vollständig begründen, erklären aber im Beispiel 4.11 die zugrundeliegende Idee. Wir konzentrieren uns stattdessen auf die Ankündigung (ii), nach geeigneter Abstraktion vom physikalischen Rahmen.

Normierte Vektorräume

Wie bereits im Kapitel 1 bemerkt, besitzen viele (aber nicht alle, siehe unten) physikalische Messgrößen (in einer gewissen Approximation) eine (affin-)lineare Struktur:

Sie (oder ihre Differenzen) können untereinander addiert und mit Skalaren multipliziert werden. Bei realistischen (endlich vielen) Operationen genügen im mathematischen Bild hierzu die rationalen Zahlen als Grundkörper. Zur Quantifizierung und Verbesserung der Approximation dient ein Abstandsbegriff, und im idealisierten Grenzfall beliebig guter Approximation ist der Übergang zu den reellen Zahlen unumgänglich.

· Typisches Beispiel für eine solche Struktur ist der (beispielsweise aus der Mechanik) vertraute n -dimensionale euklidische Raum, die Menge der n -Tupel reeller Zahlen

$$\mathbb{R}^n = \left\{ x = (x^1, \dots, x^n)^T = \begin{pmatrix} x^1 \\ \vdots \\ x^n \end{pmatrix} \mid x^i \in \mathbb{R} \right\} \quad (4.1)$$

mit den üblichen Verabredungen. (Beachte insbesondere, dass wir hauptsächlich Spaltenvektoren und die Koordinatenindizes oben schreiben.) Hier ist “der Abstand” zwischen $x, y \in \mathbb{R}^n$ die nicht-negative reelle Zahl

$$d_E(x, y) := \left(\sum_{i=1}^n (x^i - y^i)^2 \right)^{1/2} \quad (4.2)$$

Zur Isolierung der wichtigen mathematischen Strukturen halten wir fest:

· Häufig ist der physikalische Konfigurationsraum der $\mathbb{R}^n$ nicht als Vektorraum, sondern nur als *affiner Raum*: Ohne Beobachter ist kein Ursprung ausgezeichnet, und eines der Merkmale des euklidischen Abstandes ist ja gerade seine Translationsinvarianz, d.h. $d(x + \xi, y + \xi) = d(x, y) \forall \xi \in \mathbb{R}^n$.

Definition 4.1. Ein affiner Raum über einem Körper $\mathbb{K}$ ist eine Menge A zusammen mit einem $\mathbb{K}$ -Vektorraum V und einer Abbildung $\delta : A \times A \rightarrow V$ mit den Eigenschaften, dass

- (i) $\delta(x, y) + \delta(y, z) = \delta(x, z) \forall x, y, z \in A$
- (ii) $\forall x \in A, v \in V \exists_1 y \in A : \delta(y, x) = v$

· Die Bedingung (ii) besagt, dass für jedes $x \in A$ die Abbildung $\delta_x : A \rightarrow V, y \mapsto \delta(y, x)$ eine Bijektion von A auf V ist. (i) drückt aus, dass diese Identifikationen (δ_x für variables $x \in A$) mit der Vektorraumstruktur auf V verträglich sind: Die Gleichung

$$\delta_z(y) = \delta(y, z) = \delta(y, x) + \delta(x, z) = \delta_x(y) + \delta(x, z) \quad (4.3)$$

besagt gerade, dass die in z basierte Identifikation, δ_z , sich von der in x basierten, δ_x durch die (von y unabhängige) Translation um $\delta(x, z)$ unterscheidet.

· Etwas intuitiv ist ein affiner Raum ein Vektorraum, in dem der Ursprung “vergessen” wurde und jetzt jeder Punkt einen gleichberechtigten Anspruch darauf erhebt.

· Ist beispielsweise V ein Untervektorraum eines “grösseren” $\mathbb{K}$ -Vektorraums W , und $w \in W, w \neq 0$, so ist die Menge $A := V + w \subset W = \{v + w \mid v \in V\}$ kein Untervektorraum, aber immer noch ein affiner Unterraum (zum Vektorraum V). (Solche affinen Räume sollten von der Lösung inhomogener linearer Gleichungssysteme bekannt sind.)

· Der euklidische Abstand entstammt auf dem $\mathbb{R}^n$ als “Vektorraum der Abstandsvektoren” dem euklidischen inneren Produkt

$$\langle \xi, \eta \rangle := \sum_{i=1}^n \xi^i \eta^i \quad (4.4)$$

§4. METRISCHE RÄUME

das physikalisch die Winkelmessung abbildet. (In diesem Zusammenhang ist (4.2) der n -dimensionale Satz des Pythagoras.) Es gilt nämlich:

$$d_E(x, y)^2 = \langle x - y, x - y \rangle \quad (4.5)$$

Zur Diskussion der Vollständigkeit benötigen wir diese zusätzliche Struktur aber nicht, sondern nur die Eigenschaften der Funktion $v \mapsto d_E(v, 0)$ (“Abstand vom Ursprung”) auf dem zugrundeliegenden Vektorraum:

Definition 4.2. Sei V ein reeller Vektorraum. Eine Norm auf V ist eine Funktion

$$\|\cdot\| : V \rightarrow \mathbb{R}_{\geq 0} \quad (4.6)$$

mit den Eigenschaften

$$\begin{aligned} \text{Homogenität: } & \|\lambda v\| = |\lambda| \cdot \|v\| \quad \forall \lambda \in \mathbb{R}, v \in V \\ \text{Positivität: } & \|v\| > 0 \quad \forall v \neq 0 \\ \text{Dreiecksungleichung: } & \|v + w\| \leq \|v\| + \|w\| \end{aligned} \quad (4.7)$$

Ein solches Paar $(V, \|\cdot\|)$ heisst normierter Vektorraum.

Bemerkungen. Man trifft die gleiche Verabredung für einen komplexen Vektorraum, mit $|\cdot|$ = Absolutbetrag auf $\mathbb{C}$, mit der Einbettung $\mathbb{Q} \rightarrow \mathbb{R}$ auch für einen rationalen Vektorraum.

Beispiel 4.3. Für $n \geq 1$ ist der $\mathbb{R}^n$ mit der euklidischen Norm (auch “2-Norm”)

$$\|\xi\|_2 := \left(\sum_{i=1}^n (\xi^i)^2 \right)^{1/2} = \langle \xi, \xi \rangle^{1/2} \quad (4.8)$$

ein normierter Vektorraum. (Der Fall $n = 1$ ist bereits aus 1.8 bekannt.) Zum Nachweis halten wir zunächst die charakteristischen Eigenschaften des euklidischen inneren Produktes (4.4) fest:

Lemma 4.4. Für alle $\xi, \eta, \zeta \in \mathbb{R}^n$ gilt

$$\begin{aligned} \text{Symmetrie: } & \langle \xi, \eta \rangle = \langle \eta, \xi \rangle \\ \text{(Bi-)Linearität: } & \langle \xi + \lambda \eta, \zeta \rangle = \langle \xi, \zeta \rangle + \lambda \langle \eta, \zeta \rangle \\ \text{Positiv Definitheit: } & \langle \xi, \xi \rangle > 0 \quad \text{falls } \xi \neq 0 \\ \text{Cauchy-Schwarz-Ungleichung: } & |\langle \xi, \eta \rangle| \leq \|\xi\|_2 \|\eta\|_2 \end{aligned} \quad (4.9)$$

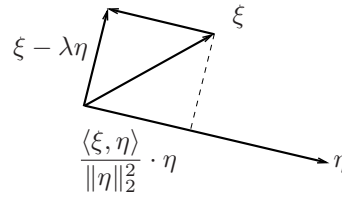
Beweis. Die ersten beiden Eigenschaften verifiziert man durch direktes Nachrechnen. Die dritte folgt aus der Positivität von Quadratzahlen (s. bereits 1.4) Zur Cauchy-Schwarz-Ungleichung: Für $\eta = 0$ ($\Leftrightarrow \|\eta\|_2 = 0$) ist die Aussage klar. Für $\eta \neq 0$, beliebiges ξ , sei

$$\lambda := \frac{\langle \xi, \eta \rangle}{\|\eta\|_2^2} \quad (4.10)$$

Dann gilt:

$$\begin{aligned}
 0 &\leq \|\xi - \lambda\eta\|_2^2 = \|\xi\|_2^2 + \lambda^2\|\eta\|_2^2 - 2\lambda\langle\xi, \eta\rangle \\
 &= \|\xi\|_2^2 + \frac{\langle\xi, \eta\rangle^2}{\|\eta\|_2^2} - 2\frac{\langle\xi, \eta\rangle^2}{\|\eta\|_2^2} \\
 &= \|\xi\|_2^2 - \frac{\langle\xi, \eta\rangle^2}{\|\eta\|_2^2} \\
 \Rightarrow \quad \langle\xi, \eta\rangle^2 &\leq \|\xi\|_2^2 \cdot \|\eta\|_2^2
 \end{aligned}
 \tag{4.11}$$

Deutung:



Durch Wurzelziehen folgt mit 1.15 die Behauptung. $\square$

Die Homogenität der euklidischen Norm (4.8) folgt nun aus der Bilinearität des inneren Produkts sowie 1.15, die Positivität aus der Definitheit. Die Dreiecksungleichung aus der Rechnung und Abschätzung

$$\|\xi + \eta\|^2 = \|\xi\|^2 + \|\eta\|^2 + 2\langle\xi, \eta\rangle \leq \|\xi\|^2 + \|\eta\|^2 + 2\|\xi\| \cdot \|\eta\| = (\|\xi\| + \|\eta\|)^2 \tag{4.12}$$

wieder durch Wurzelziehen. $\square$

Obwohl geometrisch anschaulich und wohlvertraut ist die euklidische Norm (u.a. wegen der Wurzel!) nicht für alle Zwecke am nützlichsten. Hier sind einige andere.

Beispiel 4.5. Die Funktion $\|\cdot\|_\infty$, definiert durch

$$\|x\|_\infty := \max\{|x^i| \mid i = 1, \dots, n\} \tag{4.13}$$

ist eine Norm auf $\mathbb{R}^n$. Homogenität und Positivität sind klar, die Dreiecksungleichung folgt aus 1.8 via

$$\|x + y\|_\infty = \max\{|x^i + y^i|\} \leq \max\{|x^i|\} + \max\{|y^i|\} = \|x\|_\infty + \|y\|_\infty \tag{4.14}$$

Ausführlicher: Es gilt $\max\{|x^i + y^i|\} = |x^{i_*} + y^{i_*}|$ für ein gewisses $i_* \in \{1, \dots, n\}$. Für dieses i_* gilt erstens $|x^{i_*} + y^{i_*}| \leq |x^{i_*}| + |y^{i_*}|$ wegen der gewöhnlichen Dreiecksungleichung, und weiters $|x^{i_*}| \leq |x^i| \ \forall i$, d.h. $|x^{i_*}| \leq \max\{|x^i|\}$, ebenso $|y^{i_*}| \leq \max\{|y^i|\}$, zusammen dann (4.14).

· Ebenso ist

$$\|x\|_1 := \sum_{i=1}^n |x^i| \tag{4.15}$$

eine Norm auf $\mathbb{R}^n$.

· In vielen (mathematisch abstrakten, aber auch physikalisch motivierten) Situationen ist es sogar zweckmässig, die Abstandsmessung ohne Bezug auf lineare Strukturen einzuführen.

Definition 4.6. Ein metrischer Raum (sc. über $\mathbb{R}$) ist eine Menge X zusammen mit einer Abbildung (“Abstandsfunktion”)

$$d_X = d : X \times X \rightarrow \mathbb{R}_{\geq 0} = \{x \in \mathbb{R} \mid x \geq 0\} \tag{4.16}$$

§4. METRISCHE RÄUME

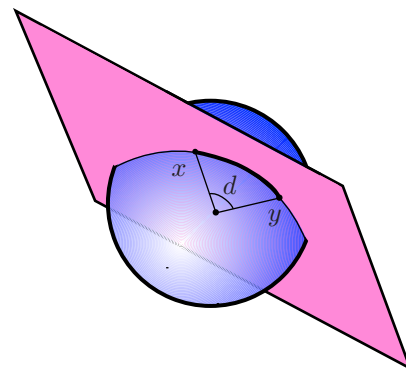
mit den Eigenschaften:

$$\begin{aligned} \text{Symmetrie: } d(x, y) &= d(y, x) \quad \forall x, y \in X \\ \text{Positivität: } d(x, y) &= 0 \Leftrightarrow x = y \\ \text{Dreiecksungleichung: } d(x, y) &\leq d(x, z) + d(z, y) \quad \forall x, y, z \in X \end{aligned} \quad (4.17)$$

Beispiel 4.7. · Auf der Einheitssphäre

$$S^2 = \{x \in \mathbb{R}^3 \mid \|x\|_2 = 1\} \subset \mathbb{R}^3 \quad (4.18)$$

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. $\forall x, y \in S^2 \langle x, y \rangle \in [-1, 1]$, 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.



· Ganz allgemein ist für irgendeine Menge X

$$d_X(x, y) = \begin{cases} 0 & \text{falls } x = y \\ 1 & \text{sonst} \end{cases} \quad (4.19)$$

eine Abstandsfunktion und damit (X, d_X) ein metrischer Raum.

· Ein praktisches Beispiel ist $\mathcal{B} =$ (Menge der Bahnhöfe der Deutschen Bahn) mit

$$d_{\mathcal{B}}(A, B) = \text{Dauer der kürzesten Verbindung von } A \text{ nach } B \quad (4.20)$$

· Zuletzt das motivierende Beispiel:

Lemma 4.8. *Sei $(V, \|\cdot\|)$ ein normierter Vektorraum. Mit $d_{\|\cdot\|}(v, w) = \|v - w\|$ als Abstandsfunktion wird (die) V (zugrundeliegende Menge) zu einem metrischen Raum. $d_{\|\cdot\|}$ erfüllt die Eigenschaften*

$$\begin{aligned} \text{Translationsinvarianz: } d_{\|\cdot\|}(v + u, w + u) &= d_{\|\cdot\|}(v, w) \\ \text{Homogenität: } d_{\|\cdot\|}(\lambda v, \lambda w) &= |\lambda| d_{\|\cdot\|}(v, w) \end{aligned} \quad (4.21)$$

· Ist umgekehrt V ein Vektorraum mit Abstandsfunktion d_V mit diesen Eigenschaften (4.21), so wird V durch $\|v\|_d := d_V(v, 0)$ zu einem normierten Vektorraum.

Beweis. Leichtes Umschreiben der drei Eigenschaften aus 4.2 auf die in 4.6 und (4.21), und umgekehrt. $\square$

Wir können nun daran gehen, den Vollständigkeitsbegriff auf metrische Räume (über $\mathbb{R}$) und damit insbesondere auf normierte Vektorräume auszudehnen. Wir erinnern hierzu kurz daran, dass wir in 3.1 Folgen (und Teilfolgen) mit Werten in

allgemeinen Mengen X eingeführt hatten, und für $X = \mathbb{R}$ an die Definition einer Cauchy-Folge in 3.16. Die wesentliche Beobachtung ist, dass im Cauchy-Kriterium 3.17 die Anordnung von $\mathbb{R}$ nur mittels des Absolutbetrags $|\cdot| : \mathbb{R} \rightarrow \mathbb{R}_{\geq 0}$ eingeht und für die meisten Eigenschaften des Konvergenzbegriffs sogar nur die Dreiecksungleichung.

Definition 4.9. · Eine Folge $(x_k)_{k \in \mathbb{N}}$ in einem metrischen Raum (X, d) heisst konvergent gegen $x \in X$ ⁵ falls $\forall \epsilon > 0$ ein K_ϵ existiert so, dass

$$d(x, x_k) < \epsilon \quad \forall k \geq K_\epsilon \quad (4.22)$$

· Eine Folge (x_k) heisst Cauchy-Folge, falls $\forall \epsilon > 0$ ein K_ϵ existiert so, dass

$$d(x_k, x_l) < \epsilon \quad \text{falls } k \geq K_\epsilon \text{ und } l \geq K_\epsilon \quad (4.23)$$

· Ein metrischer Raum heisst Cauchy-vollständig (oder nur vollständig), falls jede Cauchy-Folge in X konvergiert.

Definition 4.10. Ein normierter Vektorraum heisst vollständig, wenn der gemäss 4.8 zugehörige metrische Raum vollständig ist. Ein vollständiger normierter Vektorraum heisst auch *Banach-Raum*.

Bemerkungen. 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. 3.-1), 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. 3.6. (Der Versuch einer entsprechenden 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 4.11. 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 4.9 (siehe Theorem 3.17). Um mit der Ordnungsvollständigkeit aus 1.11 zu vergleichen, und damit der Ankündigung (i) auf S. 27 etwas Substanz zu verleihen, bemerken wir Folgendes: Ist $\mathbb{K}$ ein angeordneter Archimedischer Körper, so können wir den Absolutbetrag $|\cdot|_{\mathbb{K}}$, definiert über (1.11), nach Thm. 1.12 auffassen als Abbildung $\mathbb{K} \rightarrow \mathbb{R}_{\geq 0}$, und man prüft sofort, dass durch $d_{\mathbb{K}}(x, y) = |x - y|_{\mathbb{K}}$ eine Abstandsfunktion im Sinne von 4.6 erklärt wird. Es gilt dann:

Beh.: Ist $\mathbb{K}$ mit dieser Abstandsfunktion ein vollständiger metrischer Raum, so ist $\mathbb{K}$ ordnungsvollständig, d.h. jede nicht-leere nach oben beschränkte Menge besitzt eine kleinste obere Schranke. M.a.W. ist $\mathbb{K} = \mathbb{R}$.

⁵Elemente von X werden normalerweise als Punkte bezeichnet, Limites von Folgen aber oft wieder als “Grenzwerte”.

§4. METRISCHE RÄUME

Bew.: (Skizze) Sei $T \subset \mathbb{K}$ nicht leer und $T \leq B$ für ein $B \in \mathbb{K}$. Mit $x \in T$ setze $[a_1, b_1] = [x - 1, B]$ und definiere rekursiv für $k = 1, 2, \dots$:

$$m_k := \frac{a_k + b_k}{2}$$

$$[a_{k+1}, b_{k+1}] := \begin{cases} [a_k, m_k] & \text{falls } [m_k, b_k] \cap T = \emptyset \\ [m_k, b_k] & \text{falls } [m_k, b_k] \cap T \neq \emptyset \end{cases} \quad (4.24)$$

Die Folge $([a_k, b_k])_{k \in \mathbb{N}}$ ist dann eine Intervallschachtelung in $\mathbb{K}$ (!) (Hier ist es wichtig, dass $|b_k - a_k|_{\mathbb{K}} = 2^{1-k}|b_1 - a_1|_{\mathbb{K}}$ wegen 1.10 bereits in $\mathbb{K}$ eine Nullfolge ist, da $\mathbb{K}$ Archimedisches 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 4.12. Die Beispiele 4.19 und 4.20 aus 4.7 sind vollständige metrische Räume (Übungsaufgabe). Ebenso ist 4.18 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 $d(n, m) = \left| \frac{1}{n} - \frac{1}{m} \right|$ 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 4.13. Der Vektorraum $\mathbb{R}^n$ mit der Norm $\|\cdot\|_{\infty}$ aus 4.13 ist vollständig.

Beweis. Sei $(x_k)_{k \in \mathbb{N}} \subset \mathbb{R}^n$ eine Cauchy-Folge. 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

$$|y^i| \leq \max\{|y^i| \mid i = 1, \dots, n\} = \|y\|_{\infty} \quad (4.25)$$

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 3.17) und behaupten, dass $\lim_{k \rightarrow \infty} x_k = x = (x^1, \dots, x^n)^T$, jetzt im Sinne von 4.9. 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. (Inwiefern dies auch für physikalisch interessante nicht-lineare kontinuierliche oder diskrete Strukturen gilt, ist eine höchstinteressante Frage, die wir aber hier weder stellen noch beantworten.)

Mit Blick auf (3.4), dass wir zur Konvergenz einer Folge das “ $\epsilon > 0$ nur bis auf einen konstanten Faktor schlagen müssen”, erklären wir:

Definition 4.14. 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 $v \in V$

$$c \cdot \|v\|_\alpha \leq \|v\|_\beta \leq C \cdot \|v\|_\alpha \quad (4.26)$$

(Physikalisch entspricht etwa das Reskalieren der Norm um einen konstanten Faktor einfach einer anderen Wahl des Massstabs.) Durch Umschreiben der Ungleichungen (4.26) auf $\frac{1}{C} \cdot \|v\|_\beta \leq \|v\|_\alpha \leq \frac{1}{c} \cdot \|v\|_\beta$ sieht man, dass dies eine Äquivalenzrelation (reflexiv, symmetrisch und transitiv) auf der Menge der Normen auf V definiert. Es gilt dann:

Lemma 4.15. (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”) 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. □

• Für ein Andermal: Für metrische Räume führt das Ersetzen der Normen in (4.26) 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 (4.26) 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 4.16. 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

$$\sum_{j=1}^n |x^j|^2 \leq n \cdot \max\{|x^i|^2\} = n \|x\|_\infty^2 \quad (4.27)$$

also $\|x\|_2 \leq \sqrt{n} \|x\|_\infty$. Es gilt also (4.26) mit $\alpha = \infty$, $\beta = 2$, $c = 1$ und $C = \sqrt{n}$.

• Mit 4.15 folgt insbesondere, dass auch der euklidische Raum mit Abstandsfunktion (4.2) vollständig ist. Dies ist kein Zufall:

Theorem 4.17. Je zwei Normen auf einem endlich-dimensionalen (reellen oder komplexen) Vektorraum sind äquivalent.

Zur Vorbereitung des Beweises etwas Kontext.

§4. METRISCHE RÄUME

Topologische Grundbegriffe

Die grundlegende Beobachtung ist die folgende: Sei (X, d) ein metrischer Raum, und $T \subset X$ eine Teilmenge. Dann erfüllt die Einschränkung

$$d_T = d|_{T \times T} : T \times T \rightarrow \mathbb{R} \quad (4.28)$$

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 4.18. (i) Eine Teilmenge A eines metrischen Raumes heisst *abgeschlossen*⁶, 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.

$$B_\epsilon(x) = \{y \in X \mid d(x, y) < \epsilon\} \subset U \quad (4.29)$$

· Insbesondere sind für $x \in X$ und $R > 0$ die offenen Kugeln (vgl. (2.16)):

$$B_R(x) = \{y \in X \mid d(x, y) < R\} \quad (4.30)$$

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 $B_\epsilon(y) \subset B_R(x)$.

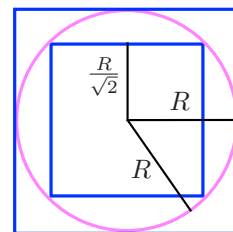
· Typische Beispiele abgeschlossener Mengen sind die “abgeschlossenen Kugeln”

$$\overline{B_R(x)} = \{y \in X \mid d(x, y) \leq R\} \quad (4.31)$$

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 (4.8) sind die $\overline{B_R}$ und B_R einfach die euklidischen Vollkugeln mit bzw. ohne die Kugelfläche. Für $d = d_\infty$ aus (4.13) sind es Quader der Kantenlänge $2R$. Die Äquivalenz der Normen 4.16 besagt, dass in jede Kugel noch ein Quader passt und umgekehrt.



· In einem metrischen Raum vom Typ (4.19) ist $\forall x \in X$:

$$B_R(x) = \begin{cases} X & \text{falls } R > 1 \\ \{x\} & \text{falls } R \leq 1 \end{cases} \quad \overline{B_R(x)} = \begin{cases} X & \text{falls } \geq 1 \\ \{x\} & \text{falls } R < 1 \end{cases} \quad (4.32)$$

· Man beachte, dass “offen” und “abgeschlossen” sich nicht gegenseitig ausschliessen. Insbesondere sind X und die leere Menge (bei geeigneter Verabredung) beides, so dass gilt:

⁶man denke sich dazu: unter Folgenkonvergenz

Lemma 4.19. $U \subset X$ ist offen $\iff X \setminus U$ 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 4.20. $U \subset X$ ist offen $\iff$ 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 4.21. (i) Ist I eine beliebige Menge, und $(U_i)_{i \in I}$ eine durch I indizierte Familie offener Mengen, so ist ihre Vereinigung

$$\bigcup_{i \in I} U_i \quad (4.33)$$

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

$$\bigcap_{i \in E} U_i \quad (4.34)$$

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 und 4.19. $\square$

§4. METRISCHE RÄUME

Der beliebige Durchschnitt offener Mengen ist nicht notwendig offen. Erkläre wo der Beweis schief geht, und gib ein Gegenbeispiel!

Lemma 4.22. *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. (4.28)). 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 4.23. Wir nennen einen metrischen Raum (X, d) beschränkt, falls die Teilmenge

$$\{d(x, y) \mid x, y \in X\} \quad (4.35)$$

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 (4.35) der Durchmesser von X , geschrieben $\text{diam } X$.

· Wir wenden diesen 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" 3.13:

Proposition 4.24. *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. $y = x$. $\square$

Zum Schluss kehren wir zu $\mathbb{R}^n$ zurück und beweisen die folgende Verallgemeinerung des Satzes von Bolzano-Weierstrass:

Proposition 4.25. *Jede beschränkte Folge $(x_k) \subset \mathbb{R}^n$ besitzt eine konvergente Teilfolge.*

Bemerkungen. 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 $(\mathbb{R}^n, \|\cdot\|_{\infty})$ (s. 4.13). Anschliessend benutzen wir dieses Resultat, um zu zeigen, dass alle Normen auf $\mathbb{R}^n$ äquivalent sind, d.h. 4.17. Daraus folgt dann mit Argumenten wie bei 4.15, dass 4.25 mit jeder Norm gilt.

Beweis von 4.25 für $\|\cdot\|_\infty$. In Imitation des Beweises von 3.14 konstruieren wir zunächst eine Folge $(A_l)_{l \in \mathbb{N}}$ wie in 4.24 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.,

$$(x_k) \subset \overline{B_R(0)} = \{y \in \mathbb{R}^n \mid \|y\|_\infty \leq R\} = [-R, R]^n \quad (4.36)$$

Laut (4.31) sind die $\overline{B_R(0)}$ (geometrisch gesehen Quader) abgeschlossen. Im ersten Schritt teilen wir $\overline{B_R(0)}$ in 2^n Teile,

$$\overline{B_{R/2}((\pm \frac{R}{2}, \pm \frac{R}{2}, \dots, \pm \frac{R}{2})^T)} \quad (4.37)$$

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}) . $\square$

Bemerkungen. Für eine unendliche Menge, ausgerüstet mit der Abstandsfunktion 4.19 ist die Aussage offenbar falsch. Wegen der Unendlichkeit von X existiert eine injektive Folge, die aber in dieser Metrik sicher keine Cauchy-Folge ist. Der Beweis scheitert daran, dass man für $R < 1$ unendlich viele Kugeln vom Radius R braucht, um X zu überdecken, vgl. (4.32).

Proposition 4.26. *Sei $\|\cdot\|$ eine beliebige Norm auf $\mathbb{R}^n$. Dann existieren Konstanten $c > 0$, $C > 0$ so, dass für alle $x \in \mathbb{R}^n$*

$$c \cdot \|x\|_\infty \leq \|x\| \leq C \cdot \|x\|_\infty \quad (4.38)$$

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 $\|\cdot\|$:

$$\|x\| \leq \sum_{i=1}^n |x^i| \cdot \|e_i\| \leq n \cdot \max\{\|e_i\|\} \cdot \|x\|_\infty \quad (4.39)$$

Damit haben wir C gefunden, und nach den Bemerkungen zu 4.15 folgt bereits, dass Konvergenz in $\|\cdot\|_\infty$ Konvergenz in $\|\cdot\|$ impliziert. (Aus (4.39) 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

$$\|x_k\| = \frac{\|\tilde{x}_k\|}{\|\tilde{x}_k\|_\infty} < \frac{1}{k} \quad \text{und} \quad \|x_k\|_\infty = 1 \quad (4.40)$$

§4. METRISCHE RÄUME

(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 4.25 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

$$\|x\| \leq \|x - x_{k_l}\| + \|x_{k_l}\| \rightarrow 0 \quad \text{für } l \rightarrow \infty \quad (4.41)$$

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, haben wir zusammengefasst sowohl 4.17 als auch 4.25 vollständig bewiesen. $\square$

Man beachte noch, dass diese Aussagen nicht für unendlich-dimensionale Vektorräume gelten. In solchen Fällen muss man die Vollständigkeit getrennt sicherstellen.

Lineare Abbildungen

Die meisten Konvergenzüberlegungen in diesem Kurs betreffen Folgen in normierten Vektorräumen. Für solche Folgen gelten Rechenregeln wie (α) und (γ) aus 3.9 sinngemäss weiter, d.h. $\lim(x_k + y_k) = \lim x_k + \lim y_k$ und $\lim\|x_k\| = \|\lim x_k\|$. Um die Verallgemeinerung der Regel (β) zu formulieren, muss auf dem Vektorraum überhaupt erst eine Multiplikation erklärt sein, und die Verträglichkeit mit der Norm schränkt die Möglichkeiten für letztere weiter ein. (Gemeint ist hier eine "echte" Multiplikation $V \times V \rightarrow V$, und nicht etwa ein "inneres Produkt" mit Werten im Grundkörper.)

• Ein Beispiel hierfür ist der Körper $\mathbb{C}$ der komplexen Zahlen. Als reeller Vektorraum ist $\mathbb{C} \cong \mathbb{R}^2$ und daher ist er wie eben gezeigt in jeder Norm als metrischer Raum über $\mathbb{R}$ vollständig. Unter allen möglichen Normen ist aber der komplexe Absolutbetrag, per Definition gleich der euklidischen Norm

$$|z| = (x^2 + y^2)^{1/2} = \|(x, y)\|_2 \quad (4.42)$$

durch die Verträglichkeit mit der Multiplikation ausgezeichnet. Für alle $z_1, z_2 \in \mathbb{C}$ gilt $|z_1 \cdot z_2| = |z_1| \cdot |z_2|$, siehe 2. Die Maximumsnorm $\|\cdot\|_\infty$ hat diese Eigenschaft nicht, denn beispielsweise ist $\|(3, 2)(2, -1)\|_\infty = 8 \neq 6 = \|(3, 2)\|_\infty \cdot \|(2, -1)\|_\infty$. Mit dem Absolutbetrag als Norm gelten dann sowohl 3.9 (β) als auch (δ) .

• Ein vergleichbares Beispiel sind die aus den Übungen bekannten Quaternionen, mit dem Unterschied, dass die quaternionische Multiplikation auf $\mathbb{R}^4$ nicht kommutativ ist.

Beispiel 4.27. Wir erinnern zunächst daran, dass für zwei (rele) Vektorräume V und W der Raum der linearen Abbildungen

$$\text{Hom}_{\mathbb{R}}(V, W) = \{A : V \rightarrow W \mid A(\lambda v_1 + v_2) = \lambda A(v_1) + A(v_2)\} \quad (4.43)$$

durch punktweise Addition und Skalarmultiplikation in natürlicher Weise zu einem reellen Vektorraum wird. $(A_1 + A_2)(v) := A_1(v) + A_2(v)$ etc.

• Ist U ein weiterer Vektorraum, so ist die multiplikativ aufgefasste Verknüpfung $\text{Hom}(W, U) \times \text{Hom}(V, W) \rightarrow \text{Hom}(V, U)$ distributiv über die Addition, d.h.

$B(A_1 + A_2) = BA_1 + BA_2$ und $(B_1 + B_2)A = B_1A + B_2A$. Im Falle $V = W$ wird damit $\text{Hom}_{\mathbb{R}}(V, V) =: \text{End}_{\mathbb{R}}(V)$ zu einer *assoziativen Algebra mit Eins* über $\mathbb{R}$. Im Unterschied zu Körpern ist die Multiplikation i.A. nicht kommutativ und es existieren nicht notwendig multiplikative Inverse.

· Sind V und W endlich-dimensional, so ist die Wahl von Basen $(e_1, \dots, e_n)$ und $(f_1, \dots, f_m)$ äquivalent zur Wahl von Isomorphismen $V \cong \mathbb{R}^n$ und $W \cong \mathbb{R}^m$ und induziert einen Isomorphismus

$$\text{Hom}_{\mathbb{R}}(V, W) \cong \text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \mathbb{R}^m) = \text{Mat}_{n \times m}(\mathbb{R}) \quad (4.44)$$

mit dem Vektorraum der $n \times m$ Matrizen mit reellen Einträgen. Wir vereinbaren im Folgenden eine solche stillschweigende Identifikation einer linearen Abbildung $A : V \rightarrow W$ mit der darstellenden Matrix

$$(A_i^j)_{\substack{j=1, \dots, m \\ i=1, \dots, n}} \in \text{Mat}_{m \times n}(\mathbb{R}) \quad (4.45)$$

· Für die Zwecke der Analysis rüsten wir nun zunächst wieder $\mathbb{R}^n$ und $\mathbb{R}^m$ mit der Maximumsnorm $\|\cdot\|_{\infty}$ aus, und stellen für $x \in \mathbb{R}^n$ fest, dass

$$\begin{aligned} \|Ax\|_{\infty} &= \max \left\{ \left| \sum_{i=1}^n A_i^j x^i \right| \mid j = 1, \dots, m \right\} \\ &\leq \max \left\{ \sum_{i=1}^n |A_i^j x^i| \mid j = 1, \dots, m \right\} \\ &\leq \max \left\{ \sum_{i=1}^n |A_i^j| \max\{|x^{\iota}| \mid \iota = 1, \dots, n\} \mid j = 1, \dots, m \right\} \\ &= \max \left\{ \sum_{i=1}^n |A_i^j| \mid j = 1, \dots, m \right\} \cdot \|x\|_{\infty} \end{aligned} \quad (4.46)$$

Die Norm von Ax lässt sich also bis auf einen von x unabhängigen Faktor durch die Norm von x abschätzen. Wegen 4.17 gilt diese Aussage dann für alle linearen Abbildungen $A : V \rightarrow W$ zwischen beliebigen endlich-dimensionalen Vektorräumen, für beliebige Wahl von Normen $\|\cdot\|_V$ und $\|\cdot\|_W$.

· Insbesondere ist die Menge

$$\{\|Ax\|_W \mid x \in V, \|x\|_V = 1\} = \left\{ \frac{\|Ax\|_W}{\|x\|_V} \mid x \in V, x \neq 0 \right\} \quad (4.47)$$

(nicht-leer, falls $n > 0$, und) nach oben beschränkt. Wir setzen ⁷

$$\|A\|_{V,W} := \sup \{ \|Ax\|_W \mid \|x\|_V = 1 \} \quad (4.48)$$

Beh.: Die Zuordnung $A \rightarrow \|A\|_{V,W}$ ist eine Norm auf dem Vektorraum $\text{Hom}_{\mathbb{R}}(V, W)$ der linearen Abbildungen, bekannt als *Matrix-* oder *Operatornorm* von A .

⁷Der Fall $n = 0$ erfordert eine getrennte Diskussion.

§4. METRISCHE RÄUME

· Anschaulich misst die Operatornorm den “maximalen Streckungsfaktor” einer linearen Abbildung. Per Definition gilt nämlich

$$\|Ax\|_W \leq \|A\|_{V,W} \cdot \|x\|_V \quad \forall x \in V \quad (4.49)$$

und keine kleinere Zahl erfüllt die gleiche Eigenschaft. (Die Aussage, dass das Maximum angenommen wird, dass also ein $x_* \in V$ existiert, mit dem in (4.49) Gleichheit gilt, lässt sich leicht mit Hilfe von 4.25 zeigen.)

Bew.: (der Eigenschaften (4.7)). Homogenität ist klar, denn

$$\{\|\lambda Ax\|_W \mid \|x\|_V = 1\} = |\lambda| \cdot \{\|Ax\|_W \mid \|x\|_V = 1\} \quad (4.50)$$

Positivität ebenfalls, denn für $A \neq 0$ existiert mindestens ein $\tilde{x} \in V$ so dass $A\tilde{x} \neq 0$. Dann ist $\tilde{x} \neq 0$, $x := \frac{\tilde{x}}{\|\tilde{x}\|_V}$ hat Norm 1 und $\|Ax\|_W > 0$. Die Menge in (4.47) enthält also ein positives Element. Zum Beweis der Dreiecksungleichung seien $A_1, A_2 \in \text{Hom}_{\mathbb{R}}(V, W)$. Dann ist wegen der Dreiecksungleichung in W $\forall x \in V$ mit $\|x\|_V = 1$

$$\begin{aligned} \|(A_1 + A_2)x\|_W &\leq \|A_1x\|_W + \|A_2x\|_W \\ &\leq \|A_1\|_W + \|A_2\|_W \end{aligned} \quad (4.51)$$

wobei wir im zweiten Schritt (4.49) verwendet haben. Daraus folgt sofort

$$\|A_1 + A_2\|_{V,W} \leq \|A_1\|_{V,W} + \|A_2\|_{V,W} \quad (4.52)$$

· Zuletzt betrachten wir die Situation mit drei normierten Vektorräumen V, W, U .

Beh.: Die Operatornorm ist *submultiplikativ*, d.h. für alle $A \in \text{Hom}_{\mathbb{R}}(V, W)$ und $B \in \text{Hom}_{\mathbb{R}}(W, U)$ gilt

$$\|BA\|_{V,U} \leq \|B\|_{W,U} \cdot \|A\|_{V,W} \quad (4.53)$$

Bew.: Durch zweimaliges Anwenden von (4.49) folgt $\forall x \in V$

$$\|BAx\|_U \leq \|B\|_{W,U} \cdot \|Ax\|_W \leq \|B\|_{W,U} \cdot \|A\|_{V,W} \cdot \|x\|_V \quad (4.54)$$

und daher gilt

$$\{\|BAx\|_U \mid \|x\|_V = 1\} \leq \|B\|_{W,U} \cdot \|A\|_{V,W} \quad (4.55)$$

und daraus folgt (4.53). $\square$

Bemerkungen. · Im Allgemeinen kann in (4.53) aber nicht = gelten, denn es gibt z.B. lineare Abbildungen $A, B \neq 0$ mit $BA = 0$.

· Die Operatornorm (4.48) hängt von der Wahl der Normen auf V und W ab. Als Normen auf dem endlich-dimensionalen reellen Vektorraum $\text{Hom}_{\mathbb{R}}(V, W)$ ($\dim = n \cdot m$) sind aber alle solche Wahlen im Sinne von 4.14 äquivalent zueinander, führen also auf den gleichen Konvergenzbegriff. Insbesondere sind sie auch äquivalent zu der Maximumsnorm

$$\|A\|_{\infty} = \max\{|A_i^j| \mid i = 1, \dots, n, j = 1, \dots, m\} \quad (4.56)$$

die aber selbst *nicht submultiplikativ* ist. (Vgl. (4.56) mit (4.46) und Gegenbeispiel).

· Für ein quantitatives Verständnis der verschiedenen Operator-Normen drängt sich im Zusammenhang mit der Interpretation (4.49) (für $V = W$) ein Vergleich mit der Eigenwerttheorie auf: Ist $\lambda \in \mathbb{R}$ ein Eigenwert von A , d.h. existiert ein $x_\lambda \neq 0$ mit $Ax_\lambda = \lambda x_\lambda$, so folgt aus der Definition (4.47) sofort, dass $\|A\|_{V,V} \geq |\lambda|$, und zwar *unabhängig von der Norm auf V* . Es gilt also auch $\|A\|_{V,V} \geq \lambda_{\max}(A) := \max\{|\lambda| \mid \lambda \text{ Eigenwert von } A\}$. Dieser maximale Eigenwert selbst ist aber keine Norm auf $\text{Hom}_{\mathbb{R}}(V, V)$, da weder positiv noch Dreiecksungleichung. Ditto für die Determinante.

· Im weiteren Verlauf dieser Vorlesung verwenden wir für den Raum der Matrizen bzw. linearen Abbildungen stets eine solche Operatornorm, in welcher $\text{Hom}_{\mathbb{R}}(V, W)$ auch vollständig ist, insbesondere für $V = W$.

· Eine assoziative Algebra mit einer submultiplikativen Norm, in welcher der zugrundeliegende Vektorraum vollständig ist, heisst auch *Banach-Algebra*. Bsp: $\text{Hom}_{\mathbb{R}}(V, V)$ für $V \neq 0$.

Eine weitere Klasse von Beispielen, die uns in dieser Vorlesung ebenfalls noch häufiger beschäftigen wird, wird durch die folgende Verallgemeinerung von $\|\cdot\|_\infty$ auf $\mathbb{R}^n$ illustriert.

Beispiel 4.28. Sei $X \neq \emptyset$ irgendeine Menge. Die Menge

$$\mathcal{B}(X, \mathbb{R}) := \{f : X \rightarrow \mathbb{R} \mid \exists M > 0 : |f(x)| \leq M \forall x \in X\} \quad (4.57)$$

der beschränkten Funktionen auf X wird durch punktweise Addition und Skalarmultiplikation zu einem (im Allgemeinen ∞ -dimensionalen) reellen Vektorraum. ($\mathbb{R}^n$ ist der Spezialfall $X = \{1, 2, \dots, n\}$.)

· Die Abbildung

$$\mathcal{B}(X, \mathbb{R}) \ni f \mapsto \|f\|_\infty = \sup\{|f(x)| \mid x \in X\} \in \mathbb{R}_{\geq 0} \quad (4.58)$$

ist eine Norm auf $\mathcal{B}(X, \mathbb{R})$ (die sup-Norm) und $(\mathcal{B}(X, \mathbb{R}), d_{\|\cdot\|_\infty})$ ist ein vollständiger metrischer Raum.

Bew.: (der Vollständigkeit) Sei $(f_k) \subset \mathcal{B}(X, \mathbb{R})$ eine Cauchy-Folge. Fixiere zunächst $x \in X$. Für jedes $\epsilon > 0$ existiert ein K_ϵ so, dass $\|f_k - f_l\|_\infty < \epsilon \forall k, l \geq K_\epsilon$. Insbesondere ist dann $|f_k(x) - f_l(x)| < \epsilon \forall k, l \geq K_\epsilon$. Daraus folgt, dass für jedes x , $(f_k(x)) \subset \mathbb{R}$ eine Cauchy-Folge ist. Wir definieren $f : X \rightarrow \mathbb{R}$ durch $f(x) = \lim f_k(x)$, und zeigen

(i) $f \in \mathcal{B}(X, \mathbb{R})$ und

(ii) $\lim \|f_k - f\|_\infty = 0$

Für $\epsilon > 0$ und K_ϵ wie eben gilt für jedes $k \geq K_\epsilon$: Für jedes x ist $|f_k(x) - f_l(x)| < \epsilon \forall l \geq K_\epsilon$. Mit $l \rightarrow \infty$ folgt $|f_k(x) - f(x)| \leq \epsilon$. Daraus folgt (i) $|f(x)| \leq |f_k(x)| + \epsilon \forall x$, d.h. f ist durch $\|f_k\|_\infty + \epsilon$ beschränkt, sowie (ii) $\|f_k - f\|_\infty < 2\epsilon$. $\square$

Der Banachsche Fixpunktsatz

Theorem/Definition 4.29. Es sei (X, d) ein metrischer Raum. Eine Abbildung $T : X \rightarrow X$ heisst *kontraktiv* (oder eine *Kontraktion*) falls eine (konstante) reelle Zahl $0 \leq c < 1$ existiert, mit der

$$d(T(x), T(y)) \leq c \cdot d(x, y) \quad \forall x, y \in X \quad (4.59)$$

§5. REIHEN

· Eine Kontraktion $T : X \rightarrow X$ auf einem nicht-leeren und vollständigen metrischen Raum besitzt genau einen Fixpunkt, d.h. $\exists_1 x_\infty \in X$ mit $T(x_\infty) = x_\infty$.

Beweis. Sei $x_0 \in X$ beliebig. Wir definieren rekursiv für $k \geq 1$ $x_k = T(x_{k-1})$.

Beh.: (x_k) ist eine Cauchy-Folge, d.h. konvergent wegen der Vollständigkeit von X .

Bew.: · Durch wiederholte Anwendung von (4.59) folgt für alle k : $d(x_k, x_{k-1}) \leq c^{k-1} \underbrace{d(x_1, x_0)}_{=:D}$

· Daraus folgt für $k > l$ mit Hilfe der Dreieckungleichung

$$\begin{aligned} d(x_k, x_l) &\leq d(x_k, x_{k-1}) + d(x_{k-1}, x_{k-2}) + \cdots + d(x_{l+1}, x_l) \leq \sum_{i=l}^{k-1} c^i \cdot D \\ &\leq c^l \cdot \frac{D}{1-c} \quad (\text{siehe geometrische Reihe im nächsten §}) \end{aligned} \quad (4.60)$$

Wegen $c < 1$ folgt $\lim c^l = 0$ und daraus sofort die Behauptung.

Beh.: Der Grenzwert $x_\infty := \lim x_k$ dieser Cauchy-Folge erfüllt $T(x_\infty) = x_\infty$.

Bew.: Für beliebiges k gilt

$$\begin{aligned} d(x_\infty, T(x_\infty)) &\leq d(x_\infty, x_k) + d(x_k, T(x_\infty)) \\ &\leq d(x_\infty, x_k) + c \cdot d(x_{k-1}, x_\infty) \rightarrow 0 \quad \text{für } k \rightarrow \infty \\ \Rightarrow d(x_\infty, T(x_\infty)) &= 0 \Rightarrow x_\infty = T(x_\infty) \end{aligned} \quad (4.61)$$

· Zum Nachweis der Eindeutigkeit sei $\tilde{x}_\infty = T(\tilde{x}_\infty)$. Aus

$$d(x_\infty, \tilde{x}_\infty) = d(T(x_\infty), T(\tilde{x}_\infty)) \leq c \cdot d(x_\infty, \tilde{x}_\infty) \quad (4.62)$$

folgt wegen $c < 1$ $d(x_\infty, \tilde{x}_\infty) = 0$. □

§5 Reihen

Begriffsklärung: ⁸ Reihen, intuitiv aufgefasst als unendliche Summen der Form

$$\sum_{k=0}^{\infty} a_k = a_0 + a_1 + a_2 + \cdots \quad (5.1)$$

sind ein wichtiger Fall der Vorstellung mathematischer Objekte durch Grenzprozesse, sowohl aus praktischer als auch aus historischer Sicht.

· Formal gesehen sind Reihen nichts anderes als Folgen $(s_n)_{n \in \mathbb{N}}$ in normierten Vektorräumen, bei denen man das Augenmerk auf die Zuwächse

$$a_k = s_k - s_{k-1} \quad (5.2)$$

anstatt auf die s_n selber richtet, m.a.W.:

⁸Dieser § ist stark an Kapitel 6 der Analysis I von Königsberger angelehnt.

Definition 5.1. Sei $(V, \|\cdot\|)$ ein normierter Vektorraum. Für eine beliebige Folge $(a_k)_{k \in \mathbb{N}_0} \subset V$ heisst die Folge $(s_n)_{n \in \mathbb{N}_0}$ mit

$$\begin{aligned} s_0 &:= a_0 \\ s_1 &:= s_0 + a_1 = a_0 + a_1 \\ s_2 &:= s_1 + a_2 = a_0 + a_1 + a_2 \\ &\vdots \\ s_n &:= s_{n-1} + a_n = \sum_{k=0}^n a_k \\ &\vdots \end{aligned} \tag{5.3}$$

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. Man schreibt

$$\sum_{k=0}^{\infty} a_k \text{ oder auch } \sum a_k \tag{5.4}$$

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) .

Lemma 5.2. Die Glieder einer konvergenten Reihe bilden eine Nullfolge (s. Def. 3.5).

Beweis. Folgt unmittelbar aus $a_k = s_k - s_{k-1}$ und den Rechenregeln für Folgen, speziell 3.9 (α). $\square$

(Die Umkehrung gilt natürlich wohlgermerkt nicht: Eine Reihe kann divergieren, auch wenn die a_k eine Nullfolge bilden. Beispiele folgen.)

· Wie angedeutet beginnt der Laufindex bei Reihen häufig bei $k = 0$. Die Konvergenz einer Reihe ändert sich aber nicht, wenn man *endlich viele* Glieder weglässt, hinzufügt (mit negativem Index), oder abändert, der Wert der Reihe aber schon.

· Unser Ziel in diesem § ist eine Übersicht über Konvergenzkriterien für Reihen sowie die Feststellung einiger Rechenregeln. Wir behandeln zunächst komplexe Reihen, mit einigen anordnungsspezifischen Bemerkungen im rein reellen Fall. Abschliessend und als Überleitung zum nächsten Kapitel richten wir den Blick auf die für die Physik wichtigen Potenzreihen.

Lemma 5.3. Sind $\sum a_k$, $\sum b_k$ zwei konvergente Reihen, und $\lambda \in \mathbb{R}$, so ist auch $\sum (a_k + \lambda b_k)$ konvergent und es gilt $\sum (a_k + \lambda b_k) = \sum a_k + \lambda \sum b_k$.

Beweis. Für die Folgen der Partialsummen gilt wegen der Assoziativität und Kommutativität bei endlichen Summen

$$\sum_{k=0}^n (a_k + \lambda b_k) = \sum_{k=0}^n a_k + \lambda \sum_{k=0}^n b_k \tag{5.5}$$

und daher folgt die Behauptung unmittelbar aus den Rechenregeln für Folgen. $\square$

§ 5. REIHEN

Konvergente Reihen bilden also, ebenso wie konvergente Folgen, mit den offensichtlichen Rechenoperationen einen (unendlich-dimensionalen) reellen 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. 5.11.

Beispiel 5.4. Für $z \in \mathbb{C}$ ist die *geometrische Reihe*

$$\sum_{k=0}^{\infty} z^k = \left(s_n = \sum_{k=0}^n z^k \right)_{n=0,1,\dots} = \left(\frac{1 - z^{n+1}}{1 - z} \right)_{n=0,1,\dots} \quad (5.6)$$

- 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 5.5. Für $s \in \mathbb{Q}$ ist die Dirichlet-Riemann-Reihe

$$\sum_{k=1}^{\infty} \frac{1}{k^s} = \left(\text{Es gibt im Allgemeinen keine geschlossene Form für die Partialsummen.} \right) \quad (5.7)$$

- 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

$$\begin{aligned} \sum_{k=1}^n \frac{1}{k^s} &\leq \sum_{k=1}^{2^l-1} \frac{1}{k^s} \\ &= 1 + \underbrace{\frac{1}{2^s} + \frac{1}{3^s}}_{\leq 2 \cdot \frac{1}{2^s}} + \underbrace{\frac{1}{4^s} + \dots + \frac{1}{7^s}}_{\leq 2^2 \cdot \frac{1}{2^s}} + \dots + \underbrace{\frac{1}{2^{(l-1)s}} + \dots + \frac{1}{(2^l-1)^s}}_{\leq 2^{l-1} \cdot \frac{1}{2^{(l-1)s}}} \\ &\leq 1 + \frac{1}{2^{s-1}} + \left(\frac{1}{2^{s-1}} \right)^2 + \dots + \left(\frac{1}{2^{s-1}} \right)^{l-1} \\ &\leq \frac{1}{1 - \frac{1}{2^{s-1}}} \quad (\text{geometrische Reihe; } 2^{1-s} < 1 \text{ für } s > 1) \end{aligned} \quad (5.8)$$

Die Folge der Partialsummen ist also monoton wachsend und nach oben beschränkt, also konvergent nach 3.12.

· Für $s \leq 1$ gilt für alle k : $\frac{1}{k^s} \geq \frac{1}{k}$. Dann gilt für $l \in \mathbb{N}$ und $n \geq 2^l$

$$\begin{aligned} \sum_{k=1}^n \frac{1}{k^s} &\geq \sum_{k=1}^{2^l} \frac{1}{k^s} \\ &= 1 + \frac{1}{2} + \underbrace{\frac{1}{3} + \frac{1}{4}}_{\geq 2 \cdot \frac{1}{4}} + \underbrace{\frac{1}{5} + \dots + \frac{1}{8}}_{\geq 4 \cdot \frac{1}{8}} + \dots + \underbrace{\frac{1}{2^{l-1}+1} + \dots + \frac{1}{2^l}}_{\geq 2^{l-1} \cdot \frac{1}{2^l}} \quad (5.9) \\ &= 1 + \frac{l}{2} \rightarrow \infty \quad \text{für } l \rightarrow \infty \end{aligned}$$

$\Rightarrow$ Die harmonische Reihe ist divergent.

Diese Beispiele illustrieren 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 5.6. 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 $n \geq m$:

$$|s_n - s_m| = \left| \sum_{k=m+1}^n a_k \right| \leq \sum_{k=m+1}^n |a_k| \leq \sum_{k=m+1}^n c_k = |u_n - u_m| < \epsilon \quad (5.10)$$

Daraus folgt mit 3.17 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 5.7. Eine Reihe $\sum a_k$ heisst absolut konvergent, falls die Reihe der Absolutbeträge $\sum |a_k|$ konvergiert.

Proposition 5.8. Eine absolut konvergente Reihe ist konvergent.

Beweis. Die Reihe der Absolutbeträge ist eine konvergente Majorante, also folgt die Aussage direkt aus 5.6. $\square$

Die Umkehrung von 5.8 gilt aber nicht, wie das folgende Beispiel zeigt:

§ 5. REIHEN

Beispiel 5.9. Die alternierende harmonische Reihe,

$$\sum_{k=1}^{\infty} \frac{(-1)^{k-1}}{k} = 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \cdots \quad (5.11)$$

konvergiert, aber nicht absolut.

Hierbei folgt die Divergenz der Reihe der Absolutbeträge aus (5.9), die Konvergenz von (5.11) aus dem sog. Leibniz-Kriterium für alternierende Reihen:

Lemma 5.10. *Sei (a_k) eine monoton fallende (oder auch wachsende) Nullfolge reeller Zahlen. Dann konvergiert*

$$\sum_{k=1}^{\infty} (-1)^k a_k \quad (5.12)$$

Beweis. Sei $s_n = \sum_{k=1}^n (-1)^k a_k$ die Folge der Partialsummen. Für $n = 2m$ gilt $s_{2m-1} \leq s_{2m}$ sowie

$$[s_{2m+1}, s_{2m+2}] \subset [s_{2m-1}, s_{2m}] \quad \text{und} \quad s_{2m} - s_{2m-1} = a_{2m} \rightarrow 0 \quad (5.13)$$

Die $([s_{2m-1}, s_{2m}])_{m=1,2,\dots}$ bilden also eine Intervallschachtelung gemäss 3.13.⁹ Es folgt sofort

$$\bigcap_{m=1}^{\infty} [s_{2m-1}, s_{2m}] = \left\{ \lim s_{2m} = \lim s_{2m-1} = \sum_{k=1}^{\infty} (-1)^k a_k \right\} \quad (5.14)$$

□

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 5.11. 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 3.3.)

Ordnen wir beispielsweise die alternierende harmonische Reihe 5.11 so um, dass auf ein positives stets zwei negative Glieder folgen, d.h.

$$1 - \frac{1}{2} - \frac{1}{4} + \frac{1}{3} - \frac{1}{6} - \frac{1}{8} + \cdots + \frac{1}{2m-1} - \frac{1}{2(2m-1)} - \frac{1}{4m} + \cdots \quad (5.15)$$

In Worten: Die m -te Gruppe enthält mit entsprechendem Vorzeichen die Kehrwerte der m -ten ungeraden Zahl, sowie der $(2m-1)$ -ten und $2m$ -ten geraden. In der

⁹Eventuell leicht verallgemeinert im Sinne von 4.24. Ein-Punkt Teilmengen von $\mathbb{R}$ werden meistens nicht als Intervalle zugelassen.

umgeordneten Reihe sind dies die Glieder Nummer $l = 3m - 2$, $3m - 1$ und $3m$. Die umordnende Bijektion ist also

$$k_l = \begin{cases} \frac{2l+1}{3} & l \equiv 1 \pmod{3} \\ 2 \frac{2l-1}{3} & l \equiv 2 \pmod{3} \\ 4 \frac{l}{3} & l \equiv 0 \pmod{3} \end{cases} \quad (5.16)$$

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 (5.15) gleich

$$\frac{1}{2} \left(1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \cdots + \frac{1}{2n-1} - \frac{1}{2n} \right) \quad (5.17)$$

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 (5.15) gegen die Hälfte von (5.11) konvergiert.

Proposition 5.12. *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 5.13. *Eine Reihe $\sum a_k$ ist genau dann absolut konvergent, wenn die Menge*

$$\mathcal{F} = \left\{ \sum_{k \in J} |a_k| \mid J \subset \mathbb{N} \text{ endliche Teilmenge} \right\} \quad (5.18)$$

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

$$\begin{array}{c} \text{Monotonie} \\ \text{von } \sum |a_k| \\ \downarrow \\ \sum_{k=1}^{\infty} |a_k| \geq \sum_{k=1}^{N_J} |a_k| \geq \sum_{k \in J} |a_k| \end{array} \quad (5.19)$$

· 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 3.12.

· 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 $\mathcal{F}$. $\square$

§ 5. REIHEN

Beweis von 5.12. Die endlichen Teilsummen der umgeordneten Reihe sind genau die gleichen wie die der ursprünglichen. Daher folgt die absolute Konvergenz der umgeordneten Reihe direkt aus 5.13. Wegen 5.8 existieren also sowohl

$$S := \sum_{k=1}^{\infty} a_k \quad \text{als auch} \quad T := \sum_{l=1}^{\infty} a_{k_l} \quad (5.20)$$

Für $\epsilon > 0$ sei nun N_ϵ so gross, dass

$$\left| S - \sum_{k=1}^{N_\epsilon} a_k \right| \leq \sum_{k=N_\epsilon+1}^{\infty} |a_k| < \epsilon \quad (5.21)$$

(Die Existenz von N_ϵ folgt aus der absoluten Konvergenz, die erste Abschätzung durch Anwenden von 5.6 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

$$\{k_1, \dots, k_{M_\epsilon}\} \supset \{1, \dots, N_\epsilon\} \quad (5.22)$$

Für alle $n \geq N_\epsilon$ und $m \geq M_\epsilon$ gilt dann

$$\left| \sum_{k=1}^n a_k - \sum_{l=1}^m a_{k_l} \right| \leq \sum_{k=N_\epsilon+1}^{\infty} |a_k| \quad (5.23)$$

Denn wegen 5.22 ist die Differenz der Partialsummen eine endliche Teilsumme der Restreihe $\sum_{k=N_\epsilon+1}^{\infty} a_k$.

Da die Partialsummen in (5.23) für $n \rightarrow \infty$ bzw. $k \rightarrow \infty$ gegen S bzw. T konvergieren, folgt mit (5.21), dass $|S - T| \leq \sum_{k=N_\epsilon+1}^{\infty} |a_k| < \epsilon$, $\forall \epsilon > 0$. $\square$

Der “grosse Umordnungssatz” ist eine Verallgemeinerung von 5.11, der Beweis sehr ähnlich.

Theorem 5.14. 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 5.3 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 5.14. (Beachte, dass die $J^{(i)}$ höchstens abzählbar unendlich sind und es höchstens abzählbar unendlich viele davon gibt, s.S. 26.)

(i) Für endliches $J^{(i)}$ sind die Aussagen trivial. Andernfalls folgt die absolute Konvergenz wie eben aus dem Lemma 5.13, und die Unabhängigkeit von der Durchnummerierung aus 5.12. Wir fixieren 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 $\forall n \geq N_\epsilon$, $h \geq H_\epsilon$, $m \geq M_\epsilon$:

$$\left| \sum_{k=1}^n a_k - \sum_{i=1}^h \sum_{l \leq m} a_{k_l^{(i)}} \right| \leq \sum_{k=N_\epsilon+1}^\infty |a_k| < \epsilon \quad (5.24)$$

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 $i = 1, \dots, h$

$$\left| S - \sum_{i=1}^h S^{(i)} \right| < \epsilon \quad (5.25)$$

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 5.14 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, wenn $\mathcal{F} = \left\{ \sum_{(k,l) \in J} |a_{kl}| \mid J \subset \mathbb{N} \times \mathbb{N} \text{ endlich} \right\}$ nach oben beschränkt ist.

• Konkret: $a_{kl} = \frac{1}{k^l}$, $k, l \geq 2$. Dann ist $\forall K, L$:

$$\sum_{k=2}^K \sum_{l=2}^L a_{kl} \leq \sum_{k=2}^K \sum_{l=2}^\infty \frac{1}{k^l} = \sum_{k=2}^K \frac{1}{k^2} \frac{1}{1 - \frac{1}{k}} \leq \sum_{k=2}^\infty \frac{1}{k(k-1)} \leq \sum_{k=1}^\infty \frac{1}{k^2} < \infty \quad (5.26)$$

sodass beliebiges Vertauschen erlaubt ist. Es folgt

$$\sum_{l=2}^\infty (\zeta(l) - 1) = \sum_{l=2}^\infty \sum_{k=2}^\infty \frac{1}{k^l} = \sum_{k=2}^\infty \sum_{l=2}^\infty \frac{1}{k^l} = \sum_{k=2}^\infty \frac{1}{k(k-1)} = 1 \quad (5.27)$$

• Beachte: Hier wird bei der Berechnung der teleskopierenden Reihe

$$\sum_{k=1}^\infty \frac{1}{k(k+1)} = \sum_{k=1}^\infty \left(\frac{1}{k} - \frac{1}{k+1} \right) \quad (5.28)$$

die Umordnung nicht über 5.14 begründet, sondern durch eine explizite Berechnung der Partialsummen:

$$\sum_{k=1}^n \left(\frac{1}{k} - \frac{1}{k+1} \right) = \sum_{k=1}^n \frac{1}{k} - \sum_{k=2}^{n+1} \frac{1}{k} = 1 - \frac{1}{n+2} \xrightarrow{n \rightarrow \infty} 1 \quad (5.29)$$

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 5.2 aber nicht hinreichend (5.9) für die Konvergenz, dass die Reihenglieder eine Nullfolge bilden. (Diese Eigenschaft ist unabhängig von der gewählten Reihenfolge.) Der Vergleich mit der geometrischen Reihe 5.4 suggeriert jedoch, dass es für die 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 der Konvergenz wiederholen wir einige Definitionen von S. 24 und 25.

-
- 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 mindestens ein Häufungswert. Ausserdem ist die Menge aller Häufungswerte beschränkt. Man nennt das Supremum der Häufungswerte den Limes Superior von (x_k) .

$$\begin{aligned} \limsup x_k &= \sup\{y \in \mathbb{R} \mid \text{Es existiert eine gegen } y \text{ konvergente Teilfolge}\} \\ &= \max\{y \in \mathbb{R} \mid \text{Es existiert eine gegen } y \text{ konvergente Teilfolge}\} \end{aligned} \quad (5.30)$$

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 5.15. *Es sei $(a_k)_{k \in \mathbb{N}}$ ein Folge komplexer Zahlen, und*

$$L := \limsup |a_k|^{1/k} \in \mathbb{R}_{\geq 0} \cup \{\infty\} \quad (5.31)$$

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

$$\sum_{k=N}^{\infty} q^k = \frac{q^N}{1-q} \quad (5.32)$$

eine konvergente Majorante für $\sum_{k=N}^{\infty} a_k$. Mit 5.6 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 5.16. Sei $r \in \mathbb{Q}$ und $z \in \mathbb{C}$. Dann ist wegen $\lim_{k \rightarrow \infty} k^{1/k} = 1$ (vgl. 3.7 (iii))

$$\limsup |k^r z^k|^{1/k} = \lim_{k \rightarrow \infty} k^{r/k} |z| = |z| \quad (5.33)$$

und daher konvergiert

$$\sum_{k=1}^{\infty} k^r z^k \quad (5.34)$$

für $|z| < 1$ absolut und divergiert für $|z| > 1$, genau wie die geometrische Reihe.

· Im Fall $r = 1$ schreiben wir die Reihe

$$\sum_{k=1}^{\infty} k z^k = z \sum_{k=0}^{\infty} (k+1) z^k = z \sum_{k=0}^{\infty} \left(\sum_{l=0}^k z^k \right) \quad (5.35)$$

als eine unendliche Summe über $k = 0, 1, \dots$ der endlichen Summen $\sum_{l=0}^k z^k$. Wegen der absoluten Konvergenz können wir über diese Indexmenge auch in einer beliebigen anderen Parametrisierung summieren, ohne das Ergebnis zu ändern (s. 5.14). Unter der Bijektion

$$\{(k, l) \mid (k, l) \in \mathbb{N}_0 \times \mathbb{N}_0, l \leq k\} \ni (k, l) \mapsto (k-l, l) =: (m, l) \in \mathbb{N}_0 \times \mathbb{N}_0 \quad (5.36)$$

(Umkehrabbildung: $k = m + l$) wird aus (5.35) insbesondere

$$\begin{aligned} z \sum_{l \leq k} z^k &= z \sum_{(l, m) \in \mathbb{N}_0 \times \mathbb{N}_0} z^{l+m} = z \sum_{l=0}^{\infty} z^l \left(\sum_{m=0}^{\infty} z^m \right) \\ &= z \left(\sum_{l=0}^{\infty} z^l \right)^2 = \frac{z}{(1-z)^2} \end{aligned} \quad (5.37)$$

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 11.3.)

Korollar 5.17. 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.

§5. REIHEN

Beweis. (Versteckt in der Voraussetzung ist die Aussage $a_k \neq 0$ für fast alle k ; $q = 0$ ist hingegen erlaubt.)

Für (i) zeigen wir $\limsup |a_k|^{1/k} \leq q$. Für $q' > q$ sei K so, dass

$$\left| \frac{a_{k+1}}{a_k} \right| < q' \quad \text{für } k \geq K \quad (5.38)$$

Dann gilt für solche k

$$\begin{aligned} |a_k| &= \left| \frac{a_k}{a_{k-1}} \right| \cdots \left| \frac{a_{K+1}}{a_K} \right| \cdot |a_K| \leq (q')^{k-K} \cdot |a_K| \\ \Rightarrow |a_k|^{1/k} &\leq q' \cdot (q'^{-K} |a_K|)^{1/k} \end{aligned} \quad (5.39)$$

Wegen 3.7(ii) folgt daraus $\limsup |a_k|^{1/k} \leq q'$. Da dies für alle $q' > q$ gilt, folgt $\limsup |a_k|^{1/k} \leq q$.

· 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 5.18. · Die Exponentialreihe $\sum_{k=0}^{\infty} \frac{z^k}{k!}$ konvergiert (für $z = 0$ sowieso und sonst) wegen

$$\frac{z^{k+1}}{(k+1)!} \cdot \frac{k!}{z^k} = \frac{z}{k+1} \xrightarrow{k \rightarrow \infty} 0 \quad (5.40)$$

für alle $z \in \mathbb{C}$ absolut.

· Die Konvergenz der hypergeometrischen Reihe

$$F(a, b, c; z) := \sum_{k=0}^{\infty} \frac{(a)_k (b)_k}{(c)_k k!} z^k \quad (5.41)$$

wird in den Übungen diskutiert.

· Man beachte, dass die Kriterien 5.15 und 5.17 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 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 5.4, 5.35, 5.40, 5.41 entsteht durch die Abhängigkeit der Reihenglieder von der komplexen Zahl, z . Konsequenterweise fasst man Reihen der Form

$$P(z) = \sum_{n=0}^{\infty} a_n z^n \quad (5.42)$$

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 $\mathbb{C}$, mehr dazu im §8.

· 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 (typischerweise: “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 5.19. *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 all n . Daraus folgt für $|z| < |z_0|$ mit $q = \left| \frac{z}{z_0} \right| < 1$:

$$|a_n z^n| = |a_n z_0^n|^n \cdot \left| \frac{z}{z_0} \right|^n \leq S q^n \quad (5.43)$$

$P(z)$ besitzt also die konvergente Majorante $\sum S q^n = \frac{S}{1-q}$ und konvergiert daher absolut. $\square$

Für gegebene Potenzreihe $P(z)$ setzen wir dann

$$R = \sup \{ r \in \mathbb{R} \mid P(r) \text{ konvergiert} \} \quad (5.44)$$

(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 $B_R(0)$ die *Konvergenzscheibe*, denn es gilt:

Theorem 5.20. *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 5.19 die absolute Konvergenz von $P(z)$.

· Wäre für ein $|z| > R$ die Reihe $P(z)$ konvergent, so wäre wegen Lemma 5.19 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 . $\square$

Durch Anwenden von Wurzel-/Quotientenkriterium 5.15, 5.17 erhalten wir dann die nützlichen Formeln zur Berechnung des Konvergenzradiuses einer Potenzreihe $\sum a_n z^n$

$$\begin{aligned} R &= \frac{1}{\limsup |a_n|^{1/n}} \\ R &= \frac{1}{\lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right|} \quad \text{falls der Grenzwert existiert} \end{aligned} \quad (5.45)$$

In diesen Formeln gilt die Vereinbarung, dass $\frac{1}{0} = \infty$ und $\frac{1}{\infty} = 0$. Beispiel: $\sum \frac{z^n}{n!}$ hat $R = \infty$, konvergiert also überall. $\sum n! z^n$ hat $R = 0$ und konvergiert nur für $z = 0$.

§ 5. REIHEN

Lemma 5.21. Für jedes $N \in \mathbb{N}$ und $r < R$ existiert eine Konstante $C \geq 0$ so, dass für alle $|z| \leq r$

$$\left| \sum_{n=0}^{\infty} a_n z^n - \sum_{n=0}^N a_n z^n \right| \leq C \cdot |z|^{N+1} \quad (5.46)$$

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. 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

$$\begin{aligned} \left| \sum_{n=N+1}^{\infty} a_n z^n \right| &\leq \sum_{n=N+1}^{\infty} |a_n z^n| = \sum_{n=N+1}^{\infty} |a_n \rho^n| \frac{|z|^n}{\rho^n} \\ &\leq S \cdot \frac{|z|^{N+1}}{\rho^{N+1}} \sum_{n=0}^{\infty} \frac{r^n}{\rho^n} = |z|^{N+1} \cdot \frac{S}{\rho^{N+1}(1 - r/\rho)} \end{aligned} \quad (5.47)$$

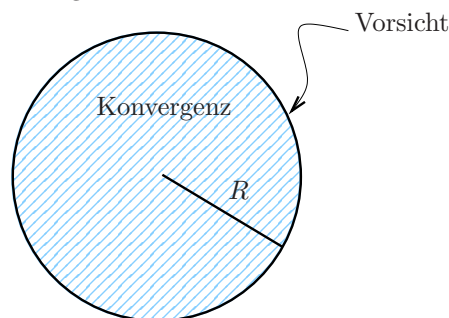
Daraus folgt die Behauptung. □

Merkwissen

I. Unter absoluter (dominierter) 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 5.6 und die Umordnungsregeln sinngemäss. Das Leibniz-Kriterium 5.10 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.¹⁰ Definiert man R wie in (5.45), so

¹⁰Wir nehmen die Existenz eines Einselements $1 \in \mathcal{A}$ an, sodass $x^0 = 1$ erklärt ist.

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

$$\sum_{n=0}^{\infty} x^n = (1 - x)^{-1} \quad \text{für } \|x\| < 1 \quad (5.48)$$

· Logarithmus-(oder Mercator-)Reihe

$$\sum_{n=1}^{\infty} (-1)^{k-1} \frac{x^k}{k} \quad R = 1 \quad (5.49)$$

· Exponentialreihe

$$\exp(x) := \sum_{n=0}^{\infty} \frac{x^n}{n!} \quad R = \infty \quad (5.50)$$

· Sinus-Reihe

$$\sin(x) := \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n+1}}{(2n+1)!} \quad R = \infty \quad (5.51)$$

· Cosinus-Reihe

$$\cos(x) := \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n}}{(2n)!} \quad R = \infty \quad (5.52)$$

· Euler-Formel (für Grundkörper $= \mathbb{C}$)

$$\exp(ix) = \cos(x) + i \sin(x) \quad (5.53)$$

Proposition 5.22. Für alle $x, y \in \mathcal{A}$ mit $xy = yx$ gilt

$$\exp(x) \exp(y) = \exp(x + y) \quad (5.54)$$

und daraus abgeleitet die bekannten Additionstheoreme der via (5.53) definierten trigonometrischen Funktionen.

Lemma 5.23. Sei $(x_n) \subset \mathcal{A}$ eine Folge mit $\lim_{n \rightarrow \infty} x_n = x$. Dann gilt

$$\lim_{n \rightarrow \infty} \left(1 + \frac{x_n}{n}\right)^n = \exp(x) \quad (5.55)$$

§5. REIHEN

Beweis von 5.22. (unter der Annahme, dass 5.23 gilt)

$$\begin{aligned}\exp(x) \exp(y) &= \lim_{n \rightarrow \infty} \left(1 + \frac{x}{n}\right)^n \lim_{n \rightarrow \infty} \left(1 + \frac{y}{n}\right)^n = (\text{Rechenregeln für Folgen}) \\ &= \lim_{n \rightarrow \infty} \left(1 + \frac{x}{n}\right)^n \left(1 + \frac{y}{n}\right)^n \quad (\text{wegen } xy = yx) \\ &= \lim_{n \rightarrow \infty} \left(1 + \frac{x}{n} + \frac{y}{n} + \frac{xy}{n^2}\right)^n\end{aligned}\tag{5.56}$$

Wegen $\lim_{n \rightarrow \infty} \left(x + y + \frac{xy}{n}\right) = x + y$ folgt die Behauptung. □

Beweis von 5.23. Für $\epsilon > 0$ sei N so gross, dass

$$\sum_{k=N+1}^{\infty} \frac{(\|x\| + 1)^k}{k!} < \epsilon \quad \text{und} \quad \|x_n\| < \|x\| + 1 \quad \forall n > N\tag{5.57}$$

Dann gilt für $n \geq N$:

$$\begin{aligned}\left\| \left(1 + \frac{x_n}{n}\right)^n - \exp(x) \right\| &\leq \underbrace{\sum_{k=0}^N \left\| \binom{n}{k} \frac{x_n^k}{n^k} - \frac{x^k}{k!} \right\|}_{\rightarrow 0 \text{ für } n \rightarrow \infty} + \underbrace{\sum_{k=N+1}^n \left\| \binom{n}{k} \frac{x_n^k}{n^k} \right\|}_{\leq \frac{(\|x\|+1)^k}{k!} < \epsilon} + \underbrace{\sum_{k=N+1}^{\infty} \frac{\|x^k\|}{k!}}_{\leq \frac{(\|x\|+1)^k}{k!} < \epsilon}\end{aligned}\tag{5.58}$$

wobei wir benutzt haben, dass

$$\binom{n}{k} \frac{1}{n^k} = \frac{1}{k!} \prod_{i=0}^{k-1} \left(1 - \frac{i}{n}\right) \begin{cases} \leq \frac{1}{k!} & \text{für den mittleren Term} \\ \xrightarrow{n \rightarrow \infty} \frac{1}{k!} & \text{für den ersten Term} \end{cases}\tag{5.59}$$

Es gibt also ein $\tilde{N} > N$ so, dass für $n \geq \tilde{N}$ der erste Term auf der rechten Seite von (5.58) $< \epsilon$ ist, alles zusammen $< 3\epsilon$. □

· Insbesondere ist $\exp(x) \exp(-x) = \exp(0) = 1$ und daher

$$\exp(x) \neq 0 \quad \forall x\tag{5.60}$$

KAPITEL 3

STETIGKEIT

Anfangen in diesem Kapitel gilt unser Hauptinteresse jetzt nicht mehr nur einzelnen reellen oder komplexen Zahlen oder Punktfolgen in metrischen Räumen, sondern vielmehr denjenigen Zusammenhängen zwischen solchen Objekten, welche durch mathematische Funktionen $F : (D, d_X|_D) \rightarrow (Y, d_Y)$ gegeben sind, nämlich als Zuordnung eines einzigen und eindeutigen Wertes $F(x)$ (in einem metrischen Raum Y) zu jedem Element x eines Definitionsbereichs D , üblicherweise Teilmenge eines metrischen Raumes X . Wir fassen solche Abbildungen dann als eigenständige mathematische Objekte auf, und untersuchen Eigenschaften der zugehörigen “Räume aller Abbildungen”. Wir wiederholen zur Klärung zunächst einige ganz wenige Grundbegriffe, die im weiteren Verlauf eine grössere Rolle spielen.

-
- Sind A, B zwei beliebige Mengen, so ist eine Abbildung von A nach B eine Teilmenge $F \subset A \times B$ mit der Eigenschaft, dass für alle $a \in A$ genau ein $F(a) \in B$ existiert so dass $(a, F(a)) \in F$. Wir denken uns dies als Zuordnung eines einzigen und eindeutigen Wertes $F(a)$ zu jedem Element $a \in A$ und schreiben $F : A \rightarrow B$.
 - Wir nennen A *Definitionsmenge* oder *Definitionsbereich*, die Menge B *Zielmenge* oder *Wertebereich* der Abbildung F . Die Menge $F(A) := \{b \in B \mid \exists a \in A : F(a) = b\} \subset B$ heisst *Bildmenge* oder *Bild* der Abbildung.
 - Der Begriff Funktion ist im Zweifel synonym mit dem Begriff Abbildung, wird im Falle $B = \mathbb{R}$ und $\mathbb{C}$ aber bevorzugt benutzt.
 - Der Definitionsbereich ist fester Bestandteil einer Abbildung. Wir schreiben $F_1 = F_2$ für zwei Abbildungen $F_1 : A_1 \rightarrow B_1$ und $F_2 : A_2 \rightarrow B_2$ genau dann wenn $A_1 = A_2$, $B_1 = B_2$ und $F_1 = F_2$ als Teilmengen von $A_1 \times B_1 = A_2 \times B_2$. (Die Gleichheit von B_1 und B_2 wird nicht immer durchgesetzt, es reicht manchmal, wenn $F_1 = F_2$ als Mengen, d.h. $A_1 = A_2$ und $F_1(a) = F_2(a) \forall a \in A_1$. Wir vereinbaren auch, dass wir Abbildungen $F : A \rightarrow B$ und $G : C \rightarrow D$ zu $G \circ F : A \rightarrow D$ verketteten können, wenn wenigstens $F(A) \subset C$.)
 - Ist $F : A \rightarrow B$ und $T \subset A$ eine Teilmenge von A , so heisst $F|_T := F \cap (T \times B) : T \rightarrow B$ die *Einschränkung* von F auf T . Im Allgemeinen und ohne weitere Annahmen kann F nicht aus $F|_T$ zurückgewonnen werden und die Aussage $F|_T = F$ ist *falsch* falls $T \subsetneq A$.
 - Eine häufige Situation ist, dass $A \subset O$ “in natürlicher Weise” Teilmenge einer grösseren Menge O ist, m.a.W., $F : A \rightarrow B$ ist “nur” auf einer Teilmenge von O definiert. Man nennt dann eine Abbildung $\bar{F} : O \rightarrow B$ mit der Eigenschaft, dass $\bar{F}|_A = F$, eine *Fortsetzung* von F nach O . Mengentheoretische Fortsetzungen existieren unter der Voraussetzung $B \neq \emptyset$ immer, wir interessieren uns aber üblicherweise für Abbildungen mit weiteren Eigenschaften. Beispiel: Fassen wir in der Abbildung von Mengen $F : \{(0, 0), (1, 2), (2, 4), (3, 5)\} : \{0, 1, 2, 3\} \rightarrow \mathbb{R}$ den Definitionsbereich in offensichtlicher Weise als Teilmenge von $\mathbb{R}$ auf, so kann F wegen $F(3) \neq 3 \cdot F(1)$ nicht zu einer linearen Abbildung von reellen Vektorräumen $\mathbb{R} \rightarrow \mathbb{R}$ fortgesetzt werden.

· Die Lösung der gewöhnlichen Differentialgleichung $\dot{x} = x^2$ mit Anfangswert $x(0) = x_0 > 0$ ist die Funktion $[0, x_0^{-1}) \rightarrow \mathbb{R}, t \mapsto x(t) = \frac{1}{x_0^{-1}-t}$ und lässt sich nicht als Lösung der DGL über $t_{\max} = x_0^{-1}$ hinaus fortsetzen.

§ 6 Stetige Abbildungen

Im Zusammenhang von metrischen Räumen ist die relevante Einschränkung an Funktionen die Grundidee der Physik, dass “kleine Änderungen der Ursachen oder Parameter kleine Folgen” haben. Am direktesten lässt sich dies mit der ersten Formulierung der Stetigkeit fassen.

Definition 6.1. Es seien $(X, d_X), (Y, d_Y)$ metrische Räume. Eine Abbildung $F : X \rightarrow Y$ *Lipschitz-stetig*, falls eine Konstante $L \geq 0$ existiert, mit der für alle $x_1, x_2 \in X$ gilt:

$$d_Y(F(x_1), F(x_2)) \leq L \cdot d_X(x_1, x_2) \quad (6.1)$$

In der Tat garantiert ja dann ein “kleiner” Abstand der Urbilder x_1 und x_2 Kontrolle über den Abstand zwischen den Bildpunkten. Je nach Wert der “Lipschitz-Konstanten” L kann es zwar vorkommen, dass $d_Y(F(x_1), F(x_2))$ noch nicht a priori “klein genug” ist, durch Verkleinern von $d_X(x_1, x_2)$ können wir aber jeden Zielwert unterbieten. Diese Beobachtung verleitet uns durch Umkehren der Beweislast zu einem allgemeineren, zunächst punktweise definierten Begriff. Wir lassen ab jetzt häufiger mal die Indices von den Abstandsfunktionen weg.

Definition 6.2. Eine Abbildung $F : X \rightarrow Y$ zwischen metrischen Räumen heisst *stetig im Punkt* $x_0 \in X$, falls für jedes $\epsilon > 0$ ein $\delta > 0$ existiert so, dass

$$d(F(x), F(x_0)) < \epsilon \quad \text{falls } d(x, x_0) < \delta \quad (6.2)$$

· Man nennt F stetig auf X , wenn F stetig in jedem $x_0 \in X$ ist.

Lemma 6.3. Eine Lipschitz-stetige Funktion ist stetig.

Beweis. Für $L = 0$ gibt es nichts zu tun. Andernfalls erfüllt, für jedes $x_0, \delta = \frac{\epsilon}{L}$ die geforderte Bedingung. $\square$

Beispiel 6.4. · Man überlegt sich, dass für eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ die ϵ - δ -Definition der Stetigkeit in $x_0 \in \mathbb{R}$ die Anschauung realisiert, dass “ f in x_0 nicht springt”: Für jedes $\epsilon > 0$ liegt für genügend kleines $\delta > 0$ der Graph der Funktion innerhalb des Rechtecks $(x_0 - \delta, x_0 + \delta) \times (f(x_0) - \epsilon, f(x_0) + \epsilon)$.

· Die Stufenfunktion

$$\mathbb{R} \ni x \mapsto \begin{cases} 1 & x \geq 0 \\ 0 & x < 0 \end{cases} \quad (6.3)$$

ist genau in $x_0 = 0$ nicht stetig.

§6. STETIGE ABBILDUNGEN

· Für $X = \mathbb{R}, \mathbb{C}$ oder eine normierte Algebra sind die algebraischen Rechenoperationen (Addition, Multiplikation, Inversenbildung) stetig. Beispielsweise gilt für gegebene $x_1, x_2 \in \mathcal{A}$ die Abschätzung

$$\begin{aligned}\|x_1^2 - x_2^2\| &= \|x_1(x_1 - x_2) + (x_1 - x_2)x_2\| \\ &\leq \|x_1\| \cdot \|x_1 - x_2\| + \|x_1 - x_2\| \cdot \|x_2\| \\ &\leq (\|x_1\| + \|x_2\|) \cdot \|x_1 - x_2\|\end{aligned}\quad (6.4)$$

Dies zeigt: $x \mapsto x^2$ ist Lipschitz-stetig in jeder Kugel $B_R(0)$. Allerdings wächst die Lipschitz-Konstante mit R und die Funktion ist daher nicht "global" Lipschitz-stetig. Insbesondere aber ist sie überall stetig.

· Polynome $P(z) = z^n + a_{n-1}z^{n-1} + \dots + a_0$ definieren stetige Funktionen: $\mathbb{C} \ni z \mapsto P(z) \in \mathbb{C}$.

· Die Wurzelfunktion $\mathbb{R}_{\geq 0} \ni x \mapsto \sqrt{x} \in \mathbb{R}_{\geq 0}$ ist stetig in $x_0 = 0$ (und natürlich auch sonst): Für gegebenes $\epsilon > 0$ ist $\sqrt{x} < \epsilon$ falls $x < \delta := \epsilon^2$, aber nicht Lipschitz-stetig: Für jedes $L > 0$ ist $\sqrt{x} > L \cdot x$ falls $0 < x < \frac{1}{L^2}$.

· Die Verkettung $X \xrightarrow{F} Y \xrightarrow{G} Z$ zweier stetiger Funktionen ist stetig: Für $x_0 \in X$, $\zeta > 0$ sei $y_0 := F(x_0) \in Y$ und $\epsilon > 0$ so, dass $d_Z(G(y), G(y_0)) < \zeta$ für $d_Y(y, y_0) < \epsilon$. Sodann sei $\delta > 0$ so, dass $d_Y(F(x), F(x_0)) < \epsilon$ für $d_X(x, x_0) < \delta$. Dann folgt $d_Z(G(F(x)), G(F(x_0))) < \zeta$ falls $d_X(x, x_0) < \delta$.

· Lineare Abbildungen zwischen endlich-dimensionalen normierten Vektorräumen über $\mathbb{R}$ oder $\mathbb{C}$ sind Lipschitz-stetig: Die Abschätzung in (4.49) besagt nichts anderes, als dass die Operatornorm eine globale Lipschitz-Konstante ist.

· Der Absolutbetrag $|\cdot| : \mathbb{C} \rightarrow \mathbb{R}_{\geq 0}$ ist die wohl wichtigste nicht-algebraische stetige Funktion: Nach der Dreiecksungleichung gilt $||z_1| - |z_2|| \leq |z_1 - z_2|$, also ist 1 eine globale Lipschitz-Konstante.

· Allgemeiner ist für jeden metrischen Raum die Abstandsfunktion $d_X : X \times X \rightarrow \mathbb{R}_{\geq 0}$ stetig als Funktion auf dem Cartesischen Produkt $X \times X$, ausgestattet mit der Produktmetrik

$$d_{X \times X}((x_1, y_1), (x_2, y_2)) := \|(d_X(x_1, x_2), d_X(y_1, y_2))\|_{\mathbb{R}^2} \quad (6.5)$$

für eine geeignete¹¹ Norm $\|\cdot\|_{\mathbb{R}^2}$ auf $\mathbb{R}^2$:

$$\begin{aligned}|d_X(x_1, y_1) - d_X(x_2, y_2)| &\leq |d_X(x_1, y_1) - d_X(x_2, y_1)| + |d_X(x_2, y_1) - d_X(x_2, y_2)| \\ (\text{Dreiecksungleichung in } X) \quad &\leq d_X(x_1, x_2) + d_X(y_1, y_2) \\ \left(\begin{array}{l} (\xi^1, \xi^2) \mapsto |\xi^1| + |\xi^2| \text{ ist eine Norm} \\ \text{auf } \mathbb{R}^2 \text{ und alle Normen auf } \mathbb{R}^2 \\ \text{sind äquivalent} \end{array} \right) &\leq C \cdot \|(d_X(x_1, x_2), d_X(y_1, y_2))\|_{\mathbb{R}^2}\end{aligned}\quad (6.6)$$

für eine geeignete Konstante C .

· Übungsaufgabe: Die Einschränkungen einer stetigen Funktion $F : X \times Y \rightarrow Z$ auf die Faktoren sind stetig.

¹¹Nicht jede Norm ist geeignet, obwohl alle Normen äquivalent sind.

· Ein exotischeres Beispiel: Die Funktion

$$\mathbb{R} \ni x \mapsto f(x) := \begin{cases} \frac{1}{q} & \text{für } x = \frac{p}{q} \in \mathbb{Q} \setminus \{0\} \text{ mit } (p, q) = 1, q > 0 \\ 1 & \text{für } x = 0 \\ 0 & \text{sonst (i.e., } x \in \mathbb{R} \setminus \mathbb{Q}) \end{cases} \quad (6.7)$$

ist stetig in jedem irrationalen Punkt $x_0 \in \mathbb{R} \setminus \mathbb{Q}$ (wähle δ so klein, dass $(x_0 - \delta, x_0 + \delta)$ keine rationalen Zahlen mit Nenner $1/q \geq \epsilon$ enthält), aber unstetig in jedem $x_0 \in \mathbb{Q}$ (für jedes $\delta > 0$ ist $x_0 + \frac{\sqrt{2}}{N} \in (x_0 - \delta, x_0 + \delta) \setminus \mathbb{Q}$ für $N > 2/\delta$).

· Die Restabschätzung 5.21 besagt, dass für eine Potenzreihe $\sum a_n x^n$ mit Konvergenzradius $R > 0$ die Zuordnung $B_R(0) \ni x \mapsto \sum a_n x^n$ in $x_0 = 0$ stetig ist. [Tatsächlich sind Potenzreihen sogar in ganz $B_R(0)$ stetig: Für jedes $\rho < R$ gilt $\forall z, w \in B_\rho(0)$:

$$\begin{aligned} \left| \sum a_n z^n - \sum a_n w^n \right| &= \left| \sum a_n (z - w)(z^{n-1} + z^{n-2}w + \dots + w^{n-1}) \right| \\ &\leq |z - w| \cdot \underbrace{\sum_{n=1}^{\infty} |a_n| n \rho^n}_{=: L_\rho} \end{aligned} \quad (6.8)$$

da die Reihe wegen $\limsup |na_n|^{1/n} = \limsup |a_n|^{1/n} = R^{-1}$ absolut konvergiert. Die Funktion $z \mapsto \sum a_n z^n$ ist also auf jedem $B_\rho(0)$ für $\rho < R$ Lipschitz-stetig, insbesondere überall stetig. Für die Exponentialreihe können wir wegen der für alle $z, w \in \mathbb{C}$ gültigen Funktionalgleichung $\exp(z + w) = \exp(z)\exp(w)$ etwas einfacher argumentieren, dass $\exp$ auf ganz $\mathbb{C}$ stetig ist: $|\exp(z_1) - \exp(z_2)| = |\exp(z_1)| \cdot |\exp(z_2 - z_1) - 1| \leq |\exp(z_1)| \cdot C \cdot |z_1 - z_2|$.

Grenzwerte von Funktionen und Häufungspunkte von Mengen

Es gibt eine sehr enge Beziehung zwischen Stetigkeit von Abbildungen und Rechenregeln für konvergente Folgen.

Proposition 6.5. *Eine Abbildung $F : X \rightarrow Y$ ist genau dann stetig in $x_0 \in X$, wenn für jede gegen x_0 konvergente Folge $(x_n)_{n \in \mathbb{N}} \subset X$ die Bildfolge $(F(x_n))_{n \in \mathbb{N}} \subset Y$ konvergiert und ihr Grenzwert*

$$\lim F(x_n) = F(x_0) = F(\lim x_n) \quad (6.9)$$

erfüllt.

Beispiel: Die Rechenregeln 3.9 sind äquivalent zur Stetigkeit der algebraischen Operationen auf $\mathbb{C}$.

Beweis. “ $\Rightarrow$ ”: Für $\epsilon > 0$ sei $\delta > 0$ so, dass $d(F(x_0), F(y)) < \epsilon$ für $d(x_0, y) < \delta$. Sodann sei N_ϵ so, dass $d(x_0, x_n) < \delta$ für $n \geq N_\epsilon$. Dann gilt für alle $n \geq N_\epsilon$: $d(F(x_0), F(x_n)) < \epsilon$.

§6. STETIGE ABBILDUNGEN

“ $\Leftarrow$ ”: Angenommen, es gäbe für ein gewisses $\epsilon_* > 0$ kein $\delta > 0$ wie benötigt. Dann gäbe es insbesondere für jedes n ein x_n mit $d(x_0, x_n) < \frac{1}{n}$, aber $d(F(x_0), F(x_n)) \geq \epsilon_*$. Dann konvergiert die Folge (x_n) gegen x_0 , aber $(F(x_n))$ nicht gegen $F(x_0)$, im Widerspruch zur Voraussetzung. $\square$

Diese Umformulierung der Stetigkeit ist einerseits praktisch zur Berechnung von Folgengrenzwerten, andererseits ein sehr nützliches Hilfsmittel zur Untersuchung stetiger Abbildungen, vor allem in Verbindung mit der Definition topologischer Grundbegriffe über Folgen wie Vollständigkeit 3.17, offene 4.18 und abgeschlossene Mengen 4.20, Beschränktheit 4.23, und andere mehr. Ein zentrales Beispiel sind die Zwänge bei der Fortsetzung stetiger Funktionen auf grössere Definitionsbereiche.

Definition 6.6. Seien X, Y metrische Räume, $D \subset X$ eine beliebige Teilmenge, $F : D \rightarrow Y$ eine Abbildung und $x_0 \in X$. Man sagt, F hat eine stetige Fortsetzung in x_0 , falls eine Abbildung $\bar{F} : D \cup \{x_0\} \rightarrow Y$ existiert, welche in x_0 stetig ist und auf $D \setminus \{x_0\}$ mit F übereinstimmt.

- Typischerweise ist $x_0 \notin D$, muss aber nicht. Häufig ist dann F auf $D \setminus \{x_0\}$ bereits stetig. Man redet natürlich auch von Fortsetzungen nach mehrpunktigen Mengen.
- Die stetige Funktion $\mathbb{R}^\times = \mathbb{R} \setminus \{0\} \ni x \mapsto x^{-1} \in \mathbb{R}$ besitzt keine stetige Fortsetzung auf ganz $\mathbb{R}$.

- Die Abbildung $f : \mathbb{Q} \rightarrow \mathbb{R}$, $f(x) := \begin{cases} 0 & \text{falls } x^2 < 2 \\ 1 & \text{falls } x^2 > 2 \end{cases}$ ist stetig auf $\mathbb{Q}$. Fassen wir

den Definitionsbereich als Teilmenge von $\mathbb{R}$ auf, so gilt: f besitzt in jedem Punkt $x \in \mathbb{R}$ ausser in $x = \sqrt{2}$ eine stetige Fortsetzung.

Lemma 6.7. F besitzt eine stetige Fortsetzung in x_0 genau dann, wenn für jede Folge $(x_n) \subset D \setminus \{x_0\}$ mit $x_n \rightarrow x_0$ die Bildfolge $F(x_n)$ in Y konvergiert.

Beweis. “ $\Rightarrow$ ”: Ist $\bar{F}$ eine stetige Fortsetzung nach x_0 , so gilt für jede Folge $(x_n) \subset D \setminus \{x_0\}$: $\bar{F}(x_n) = F(x_n) \forall n$. Konvergiert dann (x_n) gegen x_0 , so folgt $F(x_n) = \bar{F}(x_n) \rightarrow \bar{F}(x_0)$ aus 6.5.

“ $\Leftarrow$ ”: Bei der Umkehrung unterscheiden wir, ob eine gegen x_0 konvergente Folge $(x_n) \subset D \setminus \{x_0\}$ existiert oder nicht. Falls nicht, so wird F mit beliebigem $\bar{F}(x_0)$ stetig in x_0 : Fast alle Glieder einer gegen x_0 konvergente Folge in D müssen mit x_0 zusammenfallen.

Falls eine gegen x_0 konvergente Folge $(x_n) \subset D \setminus \{x_0\}$ existiert, so setzen wir $\bar{F}(x_0) := \lim F(x_n)$ für eine beliebige solche Folge.

Beh.: $\bar{F}$ ist stetig in x_0 .

Bew.: (i) Sei zunächst $(\tilde{x}_n) \subset D \setminus \{x_0\}$ eine weitere gegen x_0 konvergente Folge. Dann konvergiert die “Reissverschlussfolge” (r_n) mit $r_{2n-1} = x_n$, $r_{2n} = \tilde{x}_n$ ebenfalls gegen x_0 . Nach Voraussetzung konvergiert also die Folge $(F(r_n))$, und zwar wegen 3.10 gegen $\bar{F}(x_0)$. Dann konvergiert auch $F(\tilde{x}_n)$ gegen diesen Wert.

(ii) Sei schliesslich $(\bar{x}_n)_{n \in \mathbb{N}} \subset D$ eine beliebige gegen x_0 konvergente Folge. Gilt $\bar{x}_n = x_0$ für fast alle n , so konvergiert $\bar{F}(\bar{x}_n)$ trivialerweise gegen $\bar{F}(x_0)$. Andernfalls sei $(\tilde{x}_k)$ die Teilfolge von $(\bar{x}_n)$, bei welcher die Folgenglieder weggelassen wurden, die mit x_0 zusammenfallen. Dann konvergiert wegen (i) $\bar{F}(\tilde{x}_k) = F(\tilde{x}_k)$ gegen $\bar{F}(x_0)$.

Nach Wiedereinfügen der Glieder welche gleich x_0 sind, folgt $\bar{F}(\bar{x}_n) \rightarrow \bar{F}(x_0)$.

Mit 6.5 folgt, dass $\bar{F}$ in x_0 stetig ist. $\square$

Wir formalisieren die im Beweis gemachte Unterscheidung zur Lage von x_0 relativ zu D :

Definition 6.8. Sei $D \subset X$ eine Teilmenge eines metrischen Raums. Ein Punkt $x_0 \in X$ heisst *Häufungspunkt von D* falls eine Folge $(x_n) \subset D \setminus \{x_0\}$ existiert, welche (in X) gegen x_0 konvergiert.

Auch hier gilt wieder: Dieser Begriff ist interessant unabhängig davon, ob x_0 in D liegt oder nicht und auch unabhängig davon, ob X vollständig ist oder nicht. Ein Punkt $x_0 \in D$, welcher kein Häufungspunkt von D ist, heisst auch “isolierter Punkt” von $D \Leftrightarrow \exists \epsilon > 0 : B_\epsilon(x_0) \cap D = \{x_0\}$

Definition 6.9. Ist x_0 ein Häufungspunkt von $D \subset X$ und $F : D \rightarrow Y$ eine Abbildung mit einer stetigen Fortsetzung in x_0 , so heisst der (eindeutige!) Wert der Fortsetzung *Grenzwert von F in x_0* ,

$$\lim_{x \rightarrow x_0} F(x) = \bar{F}(x_0) \quad (6.10)$$

Ein paar weitere Beispiele: ($x \in \mathbb{R}$)

$$\lim_{x \rightarrow 1} \frac{\sqrt{x} - 1}{x - 1} = \frac{1}{2} \quad (6.11)$$

(Denn $\frac{1}{\sqrt{x}+1}$ ist eine stetige Fortsetzung in 1!)

$$\lim_{x \rightarrow 0} \frac{\sin x}{x} = 1 \quad (6.12)$$

(Benutze die Potenzreihe (5.51)!)

• Die Funktion f aus (6.7) hat den Grenzwert 0 in jedem rationalen Punkt $x \in \mathbb{Q}$ (und natürlich in jedem irrationalen, wo sie stetig ist). Eine Variante ist

$$g(x) := \begin{cases} x & \text{für } x \in \mathbb{R} \setminus \mathbb{Q} \\ 0 & \text{für } x \in \mathbb{Q} \end{cases} \quad (6.13)$$

Dann ist 0 der einzige Punkt, in dem g einen Grenzwert hat oder stetig ist.

Achtung: Existiert eine stetige Fortsetzung, so kann der Grenzwert der Funktion über den Grenzwert einer beliebigen Folge berechnet werden. Dies reicht aber nicht aus: $f : \mathbb{R}^2 \setminus \{(0,0)\} \rightarrow \mathbb{R}$,

$$f(x, y) := \frac{xy}{x^2 + y^2} \quad (6.14)$$

erfüllt $\lim_{n \rightarrow \infty} f(\frac{1}{n}, 0) = 0 = \lim_{n \rightarrow \infty} f(0, \frac{1}{n})$ aber $\lim_{n \rightarrow \infty} f(\frac{1}{n}, \frac{1}{n}) = \frac{1}{2}$ und besitzt keine stetige Fortsetzung in $(0,0)$.

• Man statte $\mathbb{N} \cup \{\infty\}$ mit einer Abstandsfunktion aus, so dass der Grenzwert von Folgen im Sinne von 3.4 mit dem Grenzwert im Sinne von 6.9 übereinstimmt (Hinweis: 4.12).

§ 7. WERTE STETIGER FUNKTIONEN

Definition 6.10. Für metrische Räume X, Y schreiben wir $\mathcal{C}(X, Y)$ für die Menge der stetigen Funktionen auf X mit Werten in Y . Sind auf Y mit der metrischen Struktur verträgliche algebraische Operationen definiert (ist z.B. $Y = V$ ein normierter Vektorraum, ausgerüstet mit der induzierten Abstandsfunktion), so übertragen sich diese Operationen auf $\mathcal{C}(X, Y)$, und es gelten relativ offensichtliche Rechenregeln für Grenzwerte. Die Idee ist einfach, dass “Verträglichkeit” der algebraischen Operationen mit der Abstandsfunktion ihre Stetigkeit impliziert. Für präzisere Formulierungen und weitere Rechenregeln für stetige Funktionen siehe Lehrbücher.

Definition 6.11. Im Falle eines angeordneten Definitionsbereichs, insbesondere für $X = \mathbb{R}$, machen die Begriffe des links- bzw. rechtsseitigen Grenzwertes Sinn: Für $D \subset \mathbb{R}$, $f : D \rightarrow Y$ und $x_0 \in \mathbb{R}$ schreiben wir (falls x_0 ein links-/rechtsseitiger Häufungspunkt von D ist, und der Grenzwert existiert) $\lim_{x \uparrow x_0} f$ für $\lim_{x \rightarrow x_0} f|_{D \cap (-\infty, x_0)}$ und $\lim_{x \downarrow x_0} f$ für $\lim_{x \rightarrow x_0} f|_{D \cap (x_0, \infty)}$.

§ 7 Werte stetiger Funktionen

Zu den wichtigsten Anwendungen der Stetigkeit von Funktionen zählt der Nachweis, dass solche Funktionen unter gewissen Voraussetzungen bestimmte Werte annehmen *müssen*. Insbesondere kann die Stetigkeit für den Beweis der Existenz von Lösungen gewisser Klassen von Problemstellungen benutzt werden. Wir diskutieren hier zunächst den Zwischenwertsatz (für stetige Funktionen mit reellen Argumenten und reellen Werten), mit einer Anwendung zur Definition der Kreiszahl π . Anschliessend stellen wir den Satz vom Maximum und Minimum vor, mit dessen Hilfe wir den Fundamentalsatz der Algebra beweisen können.

Theorem 7.1. Sei $[a, b]$ ein Intervall und $f : [a, b] \rightarrow \mathbb{R}$ stetig. Dann existiert für jedes $t \in [0, 1]$ ein $c \in [a, b]$ mit

$$f(c) = tf(a) + (1 - t)f(b) =: f_t \quad (7.1)$$

Anschaulich: Die Funktion nimmt alle Werte zwischen $f(a)$ und $f(b)$ (mindestens einmal) an.

Beweis. Für jedes feste $t \in [0, 1]$ liegt f_t “zwischen” $f(a)$ und $f(b)$. Wir betrachten den Fall $f(a) \leq f_t \leq f(b)$ (der Andere geht ist analog). Setze $[a_0, b_0] = [a, b]$ und rekursiv für $n > 0$: $m_{n-1} := \frac{a_{n-1} + b_{n-1}}{2}$ und $[a_n, b_n] = [a_{n-1}, m_{n-1}]$ oder $[m_{n-1}, b_{n-1}]$ je nachdem, ob $f(m_{n-1}) \geq f_t$ oder $f(m_{n-1}) < f_t$. Dies ergibt eine Intervallschachtelung $([a_n, b_n])$ mit der zusätzlichen Eigenschaft, dass

$$f(a_n) \leq f_t \leq f(b_n) \quad \forall n \quad (7.2)$$

Sei $c := \lim a_n = \lim b_n$ die durch diese I.S. definiert Zahl. Es gilt $c \in [a, b]$. Aus der Stetigkeit folgt $f(c) = \lim f(a_n) = \lim f(b_n)$ und daher wegen (7.2) $f(c) = f_t$ (vgl. Sandwich-Lemma 3.8). $\square$

Korollar 7.2. Sei $I \subset \mathbb{R}$ ein Intervall (geschlossen, offen, beschränkt oder unbeschränkt) und $f : I \rightarrow \mathbb{R}$ streng monoton und stetig. Dann bildet f I bijektiv auf ein Intervall J ab, und die Umkehrfunktion $f^{-1} : J \rightarrow I$ ist ebenfalls streng monoton und stetig.

Beweis. Injektivität folgt aus der streng oBdA wachsenden Monotonie: $x_1 > x_2 \Rightarrow f(x_1) > f(x_2)$. Dass $J := f(I)$ ein Intervall ist, folgt aus 7.1 (sogar unabhängig von der Monotonie). Zum Nachweis der Stetigkeit der Umkehrfunktion f^{-1} bei $y_0 = f(x_0) \in J$ zeigen wir, dass für jedes $\epsilon > 0$ ein $\delta > 0$ existiert so, dass

$$f^{-1}((y_0 - \delta, y_0 + \delta) \cap J) \subset (x_0 - \epsilon, x_0 + \epsilon) \cap I =: I_{x_0, \epsilon} \quad (7.3)$$

Ist x_0 kein Randpunkt von I , so enthält das Intervall $I_{x_0, \epsilon}$ Punkte sowohl grösser als auch kleiner als x_0 . Wegen der Monotonie enthält daher das Intervall $f(I_{x_0, \epsilon}) \subset J$ Werte sowohl kleiner als auch grösser als y_0 , und damit ein ganzes offenes Intervall $J_{y_0, \delta} := (y_0 - \delta, y_0 + \delta)$ für ein $\delta > 0$. Dieses Intervall erfüllt per Konstruktion (7.3).

Ist x_0 ein unterer Randpunkt von I , enthält also $I_{x_0, \epsilon}$ zwar Punkte grösser als x_0 , aber keine kleineren, so ist y_0 ein unterer Randpunkt von J , so dass ein $\delta > 0$ existiert mit $J_{y_0, \delta} = (y_0 - \delta, y_0 + \delta) \cap J = [y_0, y_0 + \delta) \subset f(I_{x_0, \epsilon})$.

Mit den notwendigen Änderungen für einen oberen Randpunkt gilt in jedem Fall $f^{-1}(J_{y_0, \delta}) \subset I_{x_0, \epsilon}$. $\square$

• Beispielsweise sind für $n \in \mathbb{N}$ die Wurzelfunktionen $[0, \infty) \ni x \mapsto x^{1/n} \in [0, \infty)$ als Umkehrung der Potenzfunktionen bijektiv und stetig, und durch Verkettung damit für jedes $s \in \mathbb{Q}$ auch $x \mapsto x^s$.

• Das bisher Gelernte lässt sich aber auch schon für die Rekonstruktion einiger interessanter Aussagen über die wichtigsten transzendenten Funktionen anwenden. Wir erinnern dazu an die Definition der Exponentialfunktion $\exp : \mathbb{C} \rightarrow \mathbb{C}$ durch die auf ganz $\mathbb{C}$ konvergente Potenzreihe (5.50).

Proposition/Definition 7.3. (i) Für $x \in \mathbb{R}$ ist $\exp(x) > 0$

(ii) $\exp : \mathbb{R} \rightarrow \mathbb{R}_{>0}$ ist streng monoton wachsend, surjektiv und stetig.

Die gemäss 7.2 stetige Umkehrfunktion $\ln : \mathbb{R}_{>0} \rightarrow \mathbb{R}$ heisst der natürliche Logarithmus. Er ist ebenfalls streng monoton wachsend und erfüllt die Funktionalgleichung

$$\ln x_1 x_2 = \ln x_1 + \ln x_2 \quad \text{für } x_1, x_2 > 0 \quad (7.4)$$

und erlaubt die Fortsetzung der Potenzfunktionen $x \mapsto x^s = \exp(\ln x^s) = \exp(s \ln x)$ ($s \in \mathbb{Q}$) auf reelle Exponenten. Für $x > 0$, $r \in \mathbb{R}$ setzen wir

$$x^r := \exp(r \ln x) \quad (7.5)$$

Beweis. (i) Ganz allgemein nehmen Potenzreihen mit reellen Koeffizienten bei reellem Argument reelle Werte an. $\exp(x) = (\exp(x/2))^2 > 0$ folgt damit aus (1.6) zusammen mit (5.60).

(ii) Für $x > 0$ folgt direkt aus der Potenzreihe $\exp(x) > 1 + x > 1$. Mit der Funktionalgleichung $\exp(x_1) = \exp(x_1 - x_2) \exp(x_2)$ folgt daraus die Monotonie. Mit $e := \exp(1) > 1^{12}$ folgt mittels 1.10, dass $\{e^n | n \in \mathbb{N}\}$ nicht nach oben beschränkt ist, und weil noch $\exp(-x) = 1/\exp(x)$ folgt $\exp(\mathbb{R}) = (0, \infty)$. (Stetigkeit hatten wir schon am Ende von 6.4 festgehalten.) Die Funktionalgleichung ist die Umkehrung derer für $\exp$. $\square$

¹²Wir schreiben ab jetzt auch e^x für $\exp(x)$ für jedes x .

§ 7. WERTE STETIGER FUNKTIONEN

Post quantitates exponentiales considerari debent arcus circulares eorumque sinus et cosinus, quia ex ipsis exponentialibus, quando imaginariis quantitatibus involvuntur, proveniunt.¹³ Auch diese Funktionen hatten wir ja für ganz allgemeine Argumente via Potenzreihen (5.51), (5.52) eingeführt und bemerkt, dass aus der Funktionalgleichung für die Exponentialfunktion bereits die Additionstheoreme für sin und cos folgen. Ein Version davon lautet:

$$\begin{aligned}\cos x_1 - \cos x_2 &= \left(\cos \frac{x_1+x_2}{2} \cos \frac{x_1-x_2}{2} - \sin \frac{x_1+x_2}{2} \sin \frac{x_1-x_2}{2} \right) \\ &\quad - \left(\cos \frac{x_1+x_2}{2} \cos \frac{x_1-x_2}{2} + \sin \frac{x_1+x_2}{2} \sin \frac{x_1-x_2}{2} \right) \\ &= -2 \sin \frac{x_1+x_2}{2} \sin \frac{x_1-x_2}{2}\end{aligned}\tag{7.6}$$

· Wir wollen nun durch Einschränkung auf reelle Argument die aus der Trigonometrie bekannten Eigenschaften wiederfinden. Zunächst folgt für $x \in \mathbb{R}$ aus $|\exp(ix)|^2 = \exp(ix)\exp(-ix) = 1$ mit Hilfe der Eulerschen Formel, dass

$$\cos^2 x + \sin^2 x = 1\tag{7.7}$$

Periodizität und Kreiszahl π sind aus der Reihendarstellung hingegen nicht offensichtlich, und wir müssen etwas ausholen.

· Für $x \in (0, 2]$ sind die Glieder der Sinus-Reihe $\frac{x^{2n+1}}{(2n+1)!}$ ab $n = 1$ betragsmässig streng monoton fallend. (Erstmalig nämlich: $\frac{x^3}{6} \leq x \cdot \frac{4}{6} < x$.) Die Leibniz-Betrachtungen für alternierende Reihen 5.10 implizieren daher für $0 < x \leq 2$:

$$x - \frac{x^3}{6} < \sin x < x\tag{7.8}$$

In ähnlicher Weise gilt für die gleichen x :

$$1 - \frac{x^2}{2} < \cos x < 1 - \frac{x^2}{2} + \frac{x^4}{24}\tag{7.9}$$

Proposition/Definition 7.4. *Auf dem Intervall $[0, 2]$ ist $\cos$ streng monoton fallend und besitzt genau eine Nullstelle, welche wir mit $\frac{\pi}{2}$ bezeichnen.*

Beweis. Aus (7.8) folgt $\sin x > \frac{1}{3}x > 0$ für $x \in (0, 2]$. Sind nun $x_1, x_2 \in [0, 2]$ mit $x_1 > x_2$ so gilt $\frac{x_1+x_2}{2} \in (0, 2]$ und $\frac{x_1-x_2}{2} \in (0, 2]$. Also folgt mit Hilfe von (7.6) $\cos x_1 - \cos x_2 < 0$. Ausserdem ist $\cos 0 = 1$ und wegen (7.9) gilt $\cos 2 < -\frac{1}{3}$. Mit dem Zwischenertsatz 7.1¹⁴ folgt die Existenz genau einer Nullstelle. $\square$

· Alles Übrige ergibt sich durch geschicktes Kombinieren von Additionstheoremen und (7.7). (Übungsaufgaben?) Insbesondere verifizieren wir damit die um (2.11) herum benutzen Eigenschaften und können somit Proposition 2.2 nun voll vertrauen, was wir bald für den Beweis des Fundamentalsatzes der Algebra benutzen werden.

¹³Euler, Introductio in analysin infinitorum, Cap. VIII

¹⁴Stetigkeit folgt natürlich aus dem Zusammenhang mit der Exponentialfunktion.

Satz vom Maximum und Minimum

Die andere Zutat zu diesem Beweis wollen wir aber in ihren natürlichen allgemeinen Kontext stellen.

Definition 7.5. Ein metrischer Raum X heisst *folgenkompakt*, falls jede Folge $(x_n) \subset X$ eine (in X !) konvergente Teilfolge besitzt.

Beispiel: Beschränkte und abgeschlossene Teilmengen des $\mathbb{R}^n$ sind nach 4.25 folgenkompakt. Nicht-leere offene Teilmengen von $\mathbb{R}^n$ sind nicht folgenkompakt.

Theorem 7.6. Sei X ein folgenkompakter metrischer Raum, und $f : X \rightarrow \mathbb{R}$ stetig. Dann ist

$$f(X) := \{f(x) \mid x \in X\} \quad (7.10)$$

beschränkt und (falls $X \neq \emptyset$) existiert ein $x_M \in X$ mit $f(x_M) = M := \sup f(X)$, und ein $x_m \in X$ mit $f(x_m) = m := \inf f(X)$.

Beweis. Wir zeigen, dass $f(X)$ nach oben beschränkt ist durch Widerspruch, und konstruieren gleichzeitig eine Maximumsstelle:

- (i) Falls $f(X)$ nicht nach oben beschränkt ist, so existiert für jedes n ein $x_n \in X$ mit $f(x_n) > n$.
- (ii) Falls $f(X)$ nach oben beschränkt ist, dann existiert für jedes n ein x_n mit $f(x_n) > M - \frac{1}{n}$.

Da X folgenkompakt ist, besitzt $(x_n) \subset X$ eine konvergente Teilfolge $(x_{n_k})_{k \in \mathbb{N}}$ mit Grenzwert $x_M \in X$. Wegen der Stetigkeit von f gilt gemäss 6.5

$$\lim_{k \rightarrow \infty} f(x_{n_k}) = f(x_M) \quad (7.11)$$

Klarerweise ist dies nicht mit (i) verträglich, also ist $f(X)$ nach oben beschränkt. Dann aber impliziert (ii) $f(x_M) \geq M$, d.h. wegen $f(X) \leq M$ gilt $f(x_M) = M$. $\square$

Als Korollar ergibt sich die folgende Verschärfung von 7.1: Für eine stetige Funktion $f : [a, b] \rightarrow \mathbb{R}$ ist $f([a, b]) = [\min(f), \max(f)]$.

Fundamentalsatz der Algebra

Theorem 7.7. Jedes nicht-konstante Polynom (in einer Variablen) mit komplexen Koeffizienten besitzt mindestens eine komplexe Nullstelle.

Beweis. Wir betrachten das Polynom

$$P(z) = z^n + a_{n-1}z^{n-1} + \cdots + a_1z + a_0 \quad n \in \mathbb{N}, a_i \in \mathbb{C} \quad (7.12)$$

und behaupten zunächst, dass die reellwertige und auf ganz $\mathbb{C}$ stetige Funktion $f(z) := |P(z)|$ eine globale Minimumsstelle besitzt.

§ 7. WERTE STETIGER FUNKTIONEN

Bew.: Für $|z| > 1$ und für $i = 0, \dots, n-1$ gilt $|z^i| \leq |z|^{n-1}$. Sei nun $R > 1$ und ausserdem $R > 2 \sum_{i=0}^{n-1} |a_i|$. Dann gilt für $|z| > R$:

$$|a_{n-1}z^{n-1} + \dots + a_1z + a_0| \leq |z|^{n-1} \cdot \sum_{i=0}^{n-1} |a_i| < |z|^{n-1} \frac{R}{2} < \frac{|z|^n}{2} \quad (7.13)$$

Daraus folgt, weiter für $|z| > R$:

$$f(z) = |P(z)| \geq |z^n| - |a_{n-1}z^{n-1} + \dots + a_0| > |z|^n - \frac{|z|^n}{2} = \frac{|z|^n}{2} > \frac{R^n}{2} \quad (7.14)$$

Nun ist die abgeschlossene Kreisscheibe $\overline{D_R(0)}$ als beschränkte und abgeschlossene Teilmenge von $\mathbb{C}$ folgenkompakt. Daher nimmt gemäss Thm. 7.6 die stetige Funktion $f(z)$ auf $\overline{D_R(0)}$ ein Minimum an. Es existiert also ein $z_m \in \overline{D_R(0)}$ mit $f(z) \geq f(z_m) =: m \quad \forall z \in \overline{D_R(0)}$. Wegen $|f(0)| = |a_0| < \frac{R}{2} < \frac{R^n}{2}$ ist sicher $f(z_m) < \frac{R^n}{2}$, also wegen (7.14) kleiner als jeder Wert ausserhalb der Kreisscheibe.

Also ist z_m eine globale Minimumsstelle auf $\mathbb{C}$:

$$f(z) \geq f(z_m) = m \quad \text{für alle } z \in \mathbb{C}. \quad (7.15)$$

· Nun wechseln wir die Perspektive und schauen uns $f(z)$ in der Umgebung von z_m an. Wir schreiben dazu zunächst P als Polynom in der Variablen $w = z - z_m$ um:

$$Q(w) := P(z_m + w) = b_0 + b_1w + \dots + b_{n-1}w^{n-1} + b_nw^n \quad (7.16)$$

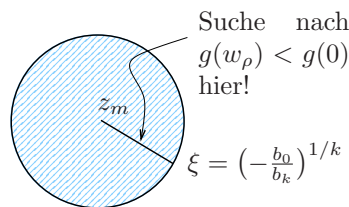
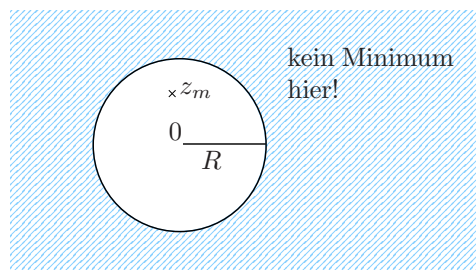
für gewisse $b_i \in \mathbb{C}$. (Tatsächlich ist $b_n = 1$, also insbesondere $\neq 0$.) Dann ist die Aussage (7.15) äquivalent dazu, dass die stetige Funktion $g(w) := |Q(w)|$ ein globales Minimum bei $w_m = 0$ besitzt. Wir behaupten, dass $m = |b_0| = 0$.

Bew.: Unter der Annahme, dass $b_0 \neq 0$, sei $k = \min\{i > 0 \mid b_i \neq 0\} \leq n$, m.a.W. ist $Q(w) = b_0 + b_k w^k + \text{höhere Potenzen}$. Gemäss 5.21 (angewandt auf eine endliche Reihe) existiert dann für $r = 1$ eine Konstante C so, dass $|Q(w) - b_0 - b_k w^k| \leq C \cdot |w|^{k+1} \quad \forall |w| \leq 1$. Gemäss 2.2 existiert ein $\xi \in \mathbb{C}^\times$ mit $\xi^k = -b_0/b_k$.

Für $0 \leq \rho < \min\{1, |\xi|^{-1}\}$ sei $w_\rho = \rho\xi$. Dann gilt $|w_\rho| < 1$ und folglich

$$g(w_\rho) = |Q(w_\rho)| \leq |b_0|(1 - \rho^k) + C \cdot \rho^{k+1} |\xi|^{k+1} = m - \rho^k(m - \rho \cdot |\xi|^{k+1} C) \quad (7.17)$$

Ist dann noch $0 < \rho < m/(|\xi|^{k+1} C)$, so folgt $g(w_\rho) < m$, im Widerspruch zu (7.15). $\square$



§ 8 Gleichmässigkeit

Der hier gegebene Beweis von 7.7 ist ein schönes Beispiel für das Zusammenspiel von Stetigkeit und Vollständigkeit zum Nachweis der *Existenz* von Lösungen eines algebraischen Problems. Er ist “nicht-konstruktiv”, indem kein explizites Verfahren zur Approximation der Lösungen gegeben wird (das aber ohne grosse Schwierigkeiten nachgereicht werden könnte). Ähnliche Prinzipien spielen auch bei der Suche nach speziellen *Funktionen*, wie etwa Lösungen von Differentialgleichungen im §15, eine wichtige Rolle.

· Wir hatten bereits beobachtet (zuletzt in 6.10), dass die Menge der Funktionen $f : X \rightarrow \mathbb{C}$ auf einer beliebigen Menge X in natürlicher Weise einen komplexen Vektorraum bildet. Ausserdem hatten wir gesehen, siehe 4.28, dass der Vektorraum der beschränkten Funktionen,

$$\mathcal{B}(X, \mathbb{C}) = \{f : X \rightarrow \mathbb{C} \mid \exists M \geq 0 : |f(x)| \leq M \forall x \in X\} \quad (8.1)$$

mit der Supremumsnorm

$$\mathcal{B}(X, \mathbb{C}) \ni f \mapsto \|f\|_\infty := \sup\{|f(x)| \mid x \in X\} \quad (8.2)$$

vollständig ist: Ist $(f_n) \subset \mathcal{B}(X, \mathbb{C})$ eine Cauchy-Folge, so ist für jedes $x \in X$ $(f_n(x)) \subset \mathbb{C}$ eine Cauchy-Folge und man setzt

$$f(x) := \lim_{n \rightarrow \infty} f_n(x) \quad (8.3)$$

Ist dann für $\epsilon > 0$ N_ϵ so, dass $\|f_n - f_m\|_\infty < \epsilon$ falls $n, m \geq N_\epsilon$. Dann gilt $\forall x \in X$, $\forall n, m \geq N_\epsilon$: $|f_n(x) - f_m(x)| < \epsilon$. Mit $m \rightarrow \infty$ folgt $|f_n(x) - f(x)| \leq \epsilon \forall x \in X$, $\forall n \geq N_\epsilon$, und daraus $f \in \mathcal{B}(X, \mathbb{C})$ und $\|f - f_n\|_\infty \rightarrow 0$ für $n \rightarrow \infty$.

· Ist X ein metrischer Raum, so ist die Menge $\mathcal{C}(X, \mathbb{C})$ der stetigen Funktionen auf X wie in 6.10 festgehalten ebenfalls in natürlicher Weise ein komplexer Vektorraum. Fassen wir $\mathcal{BC}(X, \mathbb{C}) := \mathcal{C}(X, \mathbb{C}) \cap \mathcal{B}(X, \mathbb{C})$ als Untervektorraum von $\mathcal{B}(X, \mathbb{C})$ auf, so gilt:

Theorem 8.1. $\mathcal{BC}(X, \mathbb{C})$ ist ein abgeschlossener Unterraum von $(\mathcal{B}(X, \mathbb{C}), \|\cdot\|_\infty)$, also insbesondere ein vollständiger metrischer Raum (s. 4.22).

Ein Teil des Interesses dieser Aussage ergibt sich beim Vertauschen zweier Arten von Grenzprozessen: Sind $(f_n) \subset \mathcal{C}(X, \mathbb{C})$ und $(x_k) \subset X$ konvergente Folgen, so impliziert 8.1, dass

$$\begin{aligned} \lim_{n \rightarrow \infty} \left(\lim_{k \rightarrow \infty} f_n(x_k) \right) &\stackrel{\text{Stetigkeit von } f_n}{=} \lim_{n \rightarrow \infty} f_n \left(\lim_{k \rightarrow \infty} x_k \right) = \lim_{n \rightarrow \infty} f_n(x) \stackrel{\text{Definition von } f}{=} \\ &= f(x) \stackrel{\text{Stetigkeit von } f}{=} \lim_{k \rightarrow \infty} f(x_k) \stackrel{\text{Definition von } f}{=} \lim_{k \rightarrow \infty} \left(\lim_{n \rightarrow \infty} f_n(x_k) \right) \end{aligned} \quad (8.4)$$

§ 8. GLEICHMÄSSIGKEIT

Beispiel: Die Folge stetiger und beschränkter Funktionen $f_n : [0, 1] \rightarrow \mathbb{R}$ mit

$$f_n(x) = x^n \quad (8.5)$$

erfüllt $\lim_{x \rightarrow 1} f_n(x) = 1$ für alle n , aber $\lim_{n \rightarrow \infty} f_n(x) = 0 \ \forall x \neq 0$. Es folgt

$$\lim_{x \rightarrow 1} \lim_{n \rightarrow \infty} f_n(x) = 0 \neq 1 = \lim_{n \rightarrow \infty} f_n(1) \quad (8.6)$$

Insbesondere ist die punktweise definierte Grenzfunktion gar nicht stetig, was in Anbetracht von 8.1 daran liegt, dass $\|f - f_n\|_\infty = 1 \ \forall n$.

(Ein noch(?) einfacheres Beispiel für die Probleme beim Vertauschen von Grenzprozessen ist die “Doppelfolge”

$$\mathbb{N} \times \mathbb{N} \mapsto a_{n,m} := \begin{cases} 1 & \text{für } n \geq m \\ -1 & \text{für } m > n \end{cases} \quad (8.7)$$

Klarerweise ist für jedes $m \lim_{n \rightarrow \infty} a_{n,m} = 1$, also $\lim_{m \rightarrow \infty} \lim_{n \rightarrow \infty} a_{n,m} = 1$. Auf der anderen Seite ist $\lim_{n \rightarrow \infty} \lim_{m \rightarrow \infty} a_{n,m} = -1$.)

Gleichmässige Konvergenz

Das dem Beweis von 8.1 zu Grunde liegende Prinzip ist der Begriff der *Gleichmässigkeit* von Grenzprozessen.

Definition 8.2. Eine Folge $(f_n)_{n \in \mathbb{N}}$ von Funktionen $f_n : X \rightarrow \mathbb{C}$ heisst *punktweise konvergent* gegen $f : X \rightarrow \mathbb{C}$, falls für jedes $x \in X$ die Folge $(f_n(x))_{n \in \mathbb{N}} \subset \mathbb{C}$ gegen $f(x)$ konvergiert.

Definition 8.3. Eine solche Folge heisst *gleichmässig konvergent*, falls für jedes $\epsilon > 0$ ein $N_\epsilon \in \mathbb{N}$ existiert so, dass für alle $x \in X$ $|f_n(x) - f(x)| < \epsilon$ falls $n \geq N_\epsilon$.

Bemerkungen. · Beachte den Unterschied: Bei der gleichmässigen Konvergenz müssen wir für gegebenes ϵ zuerst ein N_ϵ finden, ab welchem Index f_n an *allen Stellen* $x \in X$ nicht weiter als ϵ von f entfernt liegt. Bei der punktweise Konvergenz muss ein solcher Index erst nach Bekanntgabe von x gefunden werden.

Lemma 8.4. (i) Eine gleichmässig konvergente Folge konvergiert punktweise.
(ii) Eine Folge von beschränkten Funktionen konvergiert genau dann gleichmässig, wenn sie in der Supremums-Norm konvergiert.

Beweis. (i) ist eine leichte Übungsaufgabe, (ii) eine einfache Umformulierung der Definitionen. $\square$

Der Begriff der gleichmässigen Konvergenz ist etwas allgemeiner als die Konvergenz in der Norm: Es gibt stetige Funktionen mit “unendlicher Supremumsnorm”, m.a.W. Funktionen die in der Supremumsnorm unendlich weit voneinander entfernt liegen. Es ist dennoch natürlich, sich den Begriff der gleichmässigen Konvergenz als “von der sup-Norm” induziert vorzustellen. (Ist X folgenkompakt, so ist gemäss 7.6 jede stetige Funktion beschränkt, so dass die Begriffe echt äquivalent sind.) Hingegen kann der Begriff der punktweise Konvergenz nicht auf eine Norm oder Abstandsfunktion zurückgeführt werden.

Theorem 8.5. Sei nun X ein metrischer Raum, und $(f_n) \subset \mathcal{C}(X, \mathbb{C})$ eine Folge von stetigen Funktionen, welche gleichmässig gegen $f : X \rightarrow \mathbb{C}$ konvergiert. Dann ist f ebenfalls stetig.

Beweis. Sei $x_0 \in X$ und $\epsilon > 0$. Sei dann N so gross, dass für $n \geq N$

$$|f_n(x) - f(x)| < \epsilon \quad \forall x \in X \quad (8.8)$$

(gleichmässige Konvergenz). Sei nun $\delta > 0$ so klein, dass $|f_N(x_0) - f_N(x)| < \epsilon$ für $d_X(x, x_0) < \delta$ (Stetigkeit von f_N). Dann folgt für solche x

$$|f(x) - f(x_0)| \leq |f(x) - f_N(x)| + |f_N(x) - f_N(x_0)| + |f_N(x_0) - f(x_0)| < 3\epsilon \quad (8.9)$$

da jeder der drei Terme kleiner ist als ϵ . $\square$

Bemerkungen. · Eine wichtige Anwendung ist der Nachweis der Stetigkeit von durch Potenzreihen innerhalb ihrer Konvergenzkugel definierten Funktionen.

Beh.: Sei $P(z) = \sum_{n=0}^{\infty} a_n z^n$ eine Potenzreihe, der Einfachheit halber mit $a_n, z \in \mathbb{C}$. Sei R der Konvergenzradius von P . Dann konvergiert für jedes $r < R$ die Folge $(p_n(z) := \sum_{k=0}^n a_k z^k)_{n \in \mathbb{N}}$ auf $\overline{B_r(0)}$ gleichmässig gegen $P(z)$.

Bew.: Sei $\epsilon > 0$. Da $P(r)$ definitionsgemäss absolut konvergiert, existiert ein N so, dass

$$\sum_{k=n+1}^{\infty} |a_k| r^k < \epsilon \quad \forall n \geq N \quad (8.10)$$

Daraus folgt aber schon für alle $z \in \overline{B_r(0)}$ und alle $n \geq N$:

$$|P(z) - p_n(z)| = \left| \sum_{k=n+1}^{\infty} a_k z^k \right| \leq \sum_{k=n+1}^{\infty} |a_k| \cdot |z|^k \leq \sum_{k=n+1}^{\infty} |a_k| r^k < \epsilon \quad (8.11)$$

Da jedes p_n stetig ist, ist also wegen 8.5 P eine stetige Funktion auf der Scheibe $\overline{B_r(0)}$ für jedes $r < R$. Zusammen folgt, dass $P(z)$ auf ganz $B_R(0)$ eine stetige Funktion definiert. (Im Allgemeinen konvergiert aber p_n nicht gleichmässig gegen P auf ganz $B_R(0)$!) $\square$

Die Aussage von 8.1. folgt nun sofort durch Zusammensetzen von 8.4 und 8.5. $\square$

Beispiel 8.6. Da der Vektorraum $\mathcal{BC}(X, \mathbb{C})$ im Allgemeinen (nämlich ausser für endliches X) unendlich-dimensional ist, ist auf ihn der Satz 4.17 nicht anwendbar und tatsächlich sind auch nicht alle Normen äquivalent zueinander.

Beispiel: Auf dem Raum der auf $X = [0, 1]$ stetigen (und wegen 7.6) automatisch beschränkten) reellen Funktionen $f : [0, 1] \rightarrow \mathbb{R}$ ist die Abbildung

$$f \mapsto \|f\|_2 := \left(\int_0^1 |f(x)|^2 dx \right)^{1/2} \quad (8.12)$$

eine Norm—Beweis wie in 4.3 aus der Gültigkeit der Cauchy-Schwarz-Ungleichung

$$\left| \int_0^1 f(x)g(x)dx \right| \leq \|f\|_2 \cdot \|g\|_2 \quad (8.13)$$

§ 8. GLEICHMÄSSIGKEIT

welche man ebenfalls wie in 4.4 aus vorweggenommenen Eigenschaften des Integrals gewinnt.

· Sei nun $f_n : [0, 1] \rightarrow \mathbb{R}$ die Funktionenfolge

$$f_n := \begin{cases} n - n^4 x & x \leq \frac{1}{n^3} \\ 0 & \text{sonst} \end{cases} \quad (8.14)$$

Es gilt:

$$\int_0^{\frac{1}{n^3}} (n - n^4 x)^2 dx = \int_0^{\frac{1}{n^3}} (n^2 - 2n^5 x + n^8 x^2) dx = \frac{n^2}{n^3} - \frac{n^5}{n^6} + \frac{n^8}{3n^9} = \frac{1}{3n} \quad (8.15)$$

d.h. $\|f_n\|_2 = \sqrt{\frac{1}{3n}} \rightarrow 0$. Die Folge (f_n) konvergiert also in der 2-Norm gegen Null. Sie konvergiert zwar auch überall auf $(0, 1]$ punktweise gegen 0, die Folge $(f_n(0))$ ist aber nicht einmal beschränkt. Die Folge konvergiert damit natürlich auch nicht in der Supremumsnorm.

Wir vereinbaren, dass wir für Funktionen stets die Supremumsnorm benutzen, häufig noch in der folgenden Verallgemeinerung.

Definition 8.7. Sei X eine beliebige Menge, und $(V, \|\cdot\|)$ ein normierter Vektorraum. Dann ist die Abbildung $\mathcal{F}(X, V) \rightarrow \mathbb{R}_{\geq 0}$,

$$F \mapsto \|F\|_X := \sup\{\|F(x)\|_V \mid x \in X\} \quad (8.16)$$

eine Norm auf dem Unterraum der Funktionen mit $\|F\|_X < \infty$.

Gleichmässige Stetigkeit

Zuletzt noch eine Verschärfung des Stetigkeitsbegriffs, den wir bei der Definition des Integrals benutzen werden.

Es seien X, Y zwei metrische Räume.

Definition 8.8. Eine Funktion $F : X \rightarrow Y$ heisst *gleichmässig stetig*, falls für jedes $\epsilon > 0$ ein $\delta > 0$ existiert so, dass

$$d_Y(F(x_1), F(x_2)) < \epsilon \quad \text{falls } d_X(x_1, x_2) < \delta \quad (8.17)$$

Bemerkungen. Was ist der Unterschied von 6.2? Bei der gleichmässigen Stetigkeit müssen wir für gegebenes ϵ *zuerst* ein δ finden, für welches die Abschätzung für jedes nahe Paar $x_1, x_2 \in X$ gilt. Bei der gewöhnlichen Stetigkeit auf ganz X muss das δ erst nach der Bekanntgabe eines der beiden Punkte x_0 gefunden werden.

Beispiel: Die Funktion $\mathbb{C}^\times = \mathbb{C} \setminus \{0\} \rightarrow \mathbb{C}$, $z \mapsto 1/z$ ist (als algebraische Funktion) auf dem gesamten Definitionsbereich stetig, aber nicht gleichmässig stetig: Für $\epsilon_* = 1$ gibt es für jedes $\delta > 0$ Punkte $z_1 \in \mathbb{C}^\times$ mit $|z_1| < \min\{1, \delta\}$. Mit $z_2 = z_1/2$ gilt dann $|z_1 - z_2| = |z_1|/2 < \delta$ aber

$$\left| \frac{1}{z_1} - \frac{1}{z_2} \right| = \frac{1}{|z_1|} > 1 = \epsilon_* \quad (8.18)$$

Ist hingegen zuerst $z_0 \neq 0$ bekannt und $\epsilon > 0$, so wählen wir $0 < \delta < \min\{\frac{|z_0|}{2}, \frac{|z_0|^2}{2}\epsilon\}$. Für $|z - z_0| < \delta$ ist dann $|z| > |z_0|/2$ und daher

$$\left| \frac{1}{z} - \frac{1}{z_0} \right| = \frac{|z - z_0|}{|z| \cdot |z_0|} < \frac{2}{|z_0|^2} \cdot \delta < \epsilon \quad (8.19)$$

Für folgenkompakte Räume kann so etwas aber nicht passieren.

Theorem 8.9. *Sei X ein folgenkompakter metrischer Raum, und Y ein beliebiger metrischer Raum. Dann ist jede Funktion $F : X \rightarrow Y$ gleichmässig stetig.*

Beweis. Andernfalls gibt es ein $\epsilon_* > 0$ für welches (8.17) für kein $\delta > 0$ erfüllt ist. Insbesondere existiert für jedes $\delta_n = \frac{1}{n}$ ein Paar $x_{1,n}, x_{2,n}$ mit

$$d_X(x_{1,n}, x_{2,n}) < \frac{1}{n} \quad \text{aber} \quad d_Y(F(x_{1,n}), F(x_{2,n})) \geq \epsilon_* \quad (8.20)$$

Da X folgenkompakt ist, besitzt $(x_{1,n})$ eine konvergente Teilfolge (x_{1,n_k}) . Sei $x \in X$ ihr Grenzwert. Da F in x stetig ist, existiert ein $\delta_* > 0$ so, dass $d_Y(F(x), F(x')) < \frac{\epsilon_*}{2}$ falls $d_X(x, x') < \delta_*$. Ist dann K so, dass $d_X(x, x_{1,n_k}) < \frac{\delta_*}{2}$ und $d_X(x_{1,n_k}, x_{2,n_k}) < \frac{\delta_*}{2}$ falls $k \geq K$, dann folgt für solche k

$$\begin{aligned} d_Y(F(x), F(x_{1,n_k})) &< \frac{\epsilon_*}{2} \\ \text{und } d_X(x, x_{2,n_k}) &< \delta_* \text{ also auch } d_Y(F(x), F(x_{2,n_k})) < \frac{\epsilon_*}{2} \end{aligned} \quad (8.21)$$

Daraus folgt aber $d_Y(F(x_{1,n_k}), F(x_{2,n_k})) < \epsilon_*$, im Widerspruch zu (8.20). $\square$

KAPITEL 4

DIFFERENTIATION

Stetige Funktionen haben die charakteristische Eigenschaft, dass unter (“genügend”) kleinen Variationen ihrer Argumente ihre Werte sich nur (“beliebig”) wenig ändern. Bei Lipschitz-stetigen Funktionen nimmt diese Kontrolle die Form einer oberen *linearen* Abschätzung an. Regelmässige Zusammenhänge zwischen physikalischen Grössen sind normalerweise nicht nur stetig, sondern lassen sich vielmehr in einer Weise *linear* approximieren, dass der “relative Fehler” im Grenzfall hinreichend kleiner Änderungen beliebig klein wird. Die Differentialrechnung verbindet auf diese Weise die bisher entwickelten Ideen zu Grenzprozessen mit der aus der linearen Algebra bekannten Theorie der Vektorräume.

Wir diskutieren zunächst den Spezialfall einer reellen Veränderlichen, dessen Rechenregeln bereits vertraut sind. Die Anordnung des Definitionsbereichs (der physikalisch als “Zeit” interpretiert werden könnte) ist später grundlegend für die Lösung (Integration) gewöhnlicher Differentialgleichungen. (Noch etwas spezieller wird es, wenn auch der Wertebereich angeordnet ist, siehe z.B. 9.5.)

Wir besprechen danach die allgemeine Differentialrechnung mehrerer reeller Veränderlicher, zu deren voller Umkehrung wir allerdings erst im WS kommen, um uns für de Rest des Sommers noch etwas mit den Besonderheiten eines (zwar nicht angeordneten, dafür aber) algebraisch abgeschlossenen Definitions- und Wertebereichs (*i.e.*, Teilmengen von $\mathbb{C}$) zu beschäftigen.

§ 9 Eine reelle Veränderliche

Definition 9.1. Sei $I \subset \mathbb{R}$ ein offenes Intervall, und $(V, \|\cdot\|)$ ein vollständiger normierter Vektorraum über $\mathbb{R}$. Eine Funktion $f : I \rightarrow V$ heisst differenzierbar im Punkt $t_0 \in I$, falls der Differenzenquotient

$$I \setminus \{t_0\} \ni t \mapsto \Delta_{f,t_0}(t) := \frac{f(t) - f(t_0)}{t - t_0} \quad (9.1)$$

stetig in t_0 fortsetzbar ist. Man nennt den Grenzwert $\dot{f}(t_0) := \lim_{t \rightarrow t_0} \Delta_{f,t_0}(t)$ die Ableitung von f in t_0 .

- f heisst differenzierbar auf I , falls f in jedem Punkt von I differenzierbar ist.
- f heisst stetig differenzierbar auf I , falls f differenzierbar auf I ist, und die Abbildung

$$\dot{f} : I \rightarrow V \quad t \mapsto \dot{f}(t) \quad (9.2)$$

stetig ist.

Bemerkungen/Beispiele 9.2. · Andere Schreibweisen für $\dot{f}(t_0)$ sind $\frac{df}{dt}(t_0)$ und $f'(t_0)$.

· Eine konstante Funktion hat die Ableitung 0, für eine affin-lineare Funktion der Form $f(t) = x_0 + vt$ ($x_0, v \in V$) gilt $\dot{f}(t) = v$.

· Mit punktweiser Addition und Skalarmultiplikation wird der Raum der differenzierbaren Funktionen $\mathcal{D}(I, V)$ (bzw. der stetig differenzierbaren Funktionen, $\mathcal{C}^1(I, V)$) zu einem reellen Vektorraum (über eine Norm unterhalten wir uns später...)

· Ist auf V eine Multiplikation erklärt, so können wir Funktionen punktweise multiplizieren und es gilt die Produktregel

$$\begin{aligned} \frac{d(f \cdot g)}{dt}(t_0) &= \lim_{t \rightarrow t_0} \frac{f(t)g(t) - f(t_0)g(t) + f(t_0)g(t) - f(t_0)g(t_0)}{t - t_0} \\ &= \lim_{t \rightarrow t_0} \frac{f(t) - f(t_0)}{t - t_0} g(t) + \lim_{t \rightarrow t_0} f(t_0) \frac{g(t) - g(t_0)}{t - t_0} \\ &= \frac{df}{dt}(t_0) \cdot g(t_0) + f(t_0) \cdot \frac{dg}{dt}(t_0) \end{aligned} \quad (9.3)$$

(Daraus folgen dann die bekannten Rechenregeln für Polynome in einer reellen Variablen.)

· Ist $V \cong \mathbb{R}^m$ für ein $m \in \mathbb{N}$ (m.a.W., ist V endlich-dimensional und hat eine ausgezeichnete Basis), so können wir $f = (f^1, \dots, f^m)^T$ als m -Tupel von reellwertigen Funktionen $f^j : I \rightarrow \mathbb{R}$, $j = 1, \dots, m$ schreiben. Es gilt dann: f ist genau dann differenzierbar wenn alle f^j differenzierbar sind, und $(f(t))^j = \dot{f}^j(t)$.

· Die Ableitung der Exponentialfunktion $t \mapsto \exp(t)$ bei $t_0 = 0$ ist

$$\frac{d}{dt} \exp(t) \Big|_{t=0} = \lim_{t \rightarrow 0} \frac{\exp(t) - 1}{t} = 1 \quad (9.4)$$

(vgl. Übungen), woraus mit Hilfe der Funktionalgleichung folgt, dass $\frac{d}{dt} \exp(tx) = \exp(tx) \cdot x$. (Und daraus folgen dann die bekannten Ableitungen der trigonometrischen Funktionen.)

· Beispiel einer differenzierbaren, aber nicht stetig differenzierbaren Funktion:

$$f(t) = \begin{cases} t^2 \sin \frac{1}{t} & t \neq 0 \\ 0 & t = 0 \end{cases} \quad (9.5)$$

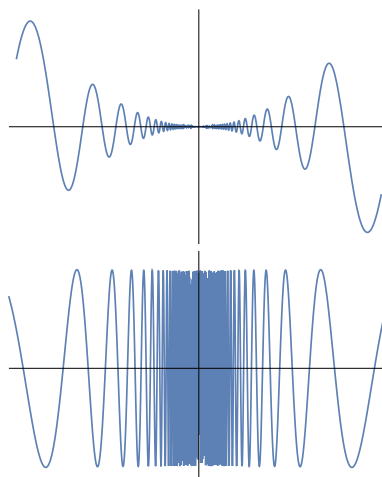
Es ist (unter Vorwegnahme der Kettenregel...)

$$\dot{f}(t) = \begin{cases} 2t \sin \frac{1}{t} - \cos \frac{1}{t} & \text{für } t \neq 0 \\ \lim_{t \rightarrow 0} t \sin \frac{1}{t} = 0 & \text{für } t = 0 \end{cases} \quad (9.6)$$

die Ableitung aber nicht stetig in $t_0 = 0$.

· Je nach Sichtweise auf Grenzwerte von Funktionen gibt es verschiedene äquivalente Formulierungen der Differenzierbarkeit. Den Einstieg in die Verallgemeinerung auf mehrere Variablen ermöglicht die folgende: f ist genau dann in t_0 differenzierbar, falls ein $\dot{f}(t_0) \in V$ existiert so, dass $\forall \epsilon > 0$ ein $\delta > 0$ existiert so, dass

$$\|f(t) - f(t_0) - \dot{f}(t_0)(t - t_0)\| < \epsilon |t - t_0| \quad \forall t \in I \cap (t_0 - \delta, t_0 + \delta) \setminus \{t_0\} \quad (9.7)$$



§9. EINE REELLE VERÄNDERLICHE

“Die Differenz zwischen Funktion und linearer Näherung geht (in der Norm) schneller als linear in $t - t_0$ gegen Null.”

Bew.: Die Bedingung (9.1) bedeutet, dass für jedes $\epsilon > 0$ ein $\delta > 0$ existiert so, dass

$$\|\Delta_{f,t_0}(t) - \dot{f}(t_0)\| < \epsilon \quad \text{falls } 0 < |t - t_0| < \delta \quad (9.8)$$

dies ist klarerweise äquivalent zu (9.7). Man kann in der Abschätzung (9.7) $<$ durch $\leq$ auf ganz $I \cap (t_0 - \delta, t_0 + \delta)$ ersetzen. $\square$

• Es ist auch sinnvoll, den Begriff auf andere Teilmengen von $\mathbb{R}$ als Definitionsbereiche zu verallgemeinern. Ist I ein abgeschlossenes oder halb-offenes Intervall, so kann man an einem Randpunkt weiter (9.7) benutzen bzw. den am Ende von §6 eingeführten Begriff des links-/rechtsseitigen Grenzwertes. Bevor man für einen allgemeinen Definitionsbereich D auf Differenzierbarkeit in $x_0 \in D$ untersucht, verlangt man, dass für hinreichend kleines δ der Durchschnitt $D \cap (x_0 - \delta, x_0 + \delta)$ ein Intervall der Form $(x_0 - \delta, x_0 + \delta)$, $[x_0, x_0 + \delta)$ oder $(x_0 - \delta, x_0]$ ist.

• Für ein kompaktes Intervall $[a, b]$ heisst eine differenzierbare Abbildung $\gamma : [a, b] \rightarrow V$ auch ein *Weg* von $x_1 = \gamma(a)$ nach $x_2 = \gamma(b)$.

• Aus (9.7) folgt sofort, dass eine differenzierbare Funktion insbesondere stetig ist:

$$\begin{aligned} \|f(t) - f(t_0)\| &\leq \|f(t) - f(t_0) - \dot{f}(t_0)(t - t_0)\| + \|\dot{f}(t_0)(t - t_0)\| \\ &\leq (\epsilon + \|\dot{f}(t_0)\|) \cdot |t - t_0| \end{aligned} \quad (9.9)$$

für $|t - t_0| < \delta$. Allgemeiner gilt:

Proposition 9.3. *Sei $f : [t_1, t_2] \rightarrow V$ eine stetige Funktion, die auf (t_1, t_2) differenzierbar ist derart, dass $\|\dot{f}\|_{(t_1, t_2)} := \sup\{\|\dot{f}(t)\| \mid t \in (t_1, t_2)\} < \infty$. (Vgl. Def. 8.7. Beispiel: f ist auf $[t_1, t_2]$ stetig differenzierbar.) Dann gilt*

$$\|f(t_2) - f(t_1)\| \leq \|\dot{f}\|_{(t_1, t_2)} \cdot |t_2 - t_1| \quad (9.10)$$

Beweis. Angenommen, die linke Seite von (9.10) wäre grösser als die rechte. Dann existierte ein $\epsilon > 0$ so, dass noch $\|f(t_2) - f(t_1)\| - (\|\dot{f}\|_{(t_1, t_2)} + \epsilon)(t_2 - t_1) > 0$. Mit solch einem ϵ betrachten wir für $t \in [t_1, t_2]$ die reellwertige Funktion

$$g_\epsilon(t) := \|f(t) - f(t_1)\| - (\|\dot{f}\|_{(t_1, t_2)} + \epsilon)(t - t_1) \quad (9.11)$$

Wegen $g_\epsilon(t_1) = 0$ ist $t_0 := \sup\{t \in [t_1, t_2] \mid g_\epsilon(t) = 0\}$ wohldefiniert. Aus der Stetigkeit von g_ϵ folgt $g_\epsilon(t_0) = 0$ und wegen $g_\epsilon(t_2) > 0$ folgt $t_0 < t_2$, und $g_\epsilon(t) > 0$ für $t \in (t_0, t_2) \neq \emptyset$. (Wäre $g_\epsilon(t') < 0$ für ein $t' \in (t_0, t_2)$ so gäbe es wegen des Zwischenwertsatzes 7.1 ein $t'' \in (t', t_2)$ mit $g_\epsilon(t'') = 0$, im Widerspruch zur Definition von t_0 .) Für alle solche t gilt:

$$\begin{aligned} 0 < g_\epsilon(t) - g_\epsilon(t_0) &= \|f(t) - f(t_1)\| - \|f(t_0) - f(t_1)\| - (\|\dot{f}\|_{(t_1, t_2)} + \epsilon)(t - t_0) \\ &\leq \|f(t) - f(t_0)\| - (\|\dot{f}\|_{(t_1, t_2)} + \epsilon)(t - t_0) \end{aligned} \quad (9.12)$$

Da aber f in t_0 differenzierbar ist, existiert ein $\delta > 0$ so, dass für solche t , die ausserdem kleiner als $t_0 + \delta$ sind, $\|f(t) - f(t_0)\| < (\|\dot{f}(t_0)\| + \epsilon)|t - t_0|$. (Siehe (9.9).) Wegen $\|\dot{f}(t_0)\| \leq \|\dot{f}\|_{(t_1, t_2)}$ steht dies im Widerspruch zu (9.12). $\square$

Korollar 9.4. *Sei $f : D \rightarrow V$ differenzierbar. Gilt $\dot{f}(t) = 0 \ \forall t \in D$, so ist f auf jedem Intervall $I \subset D$ konstant.*

Wertebereich $\mathbb{R}$

Wie bei stetigen Funktionen auch erlauben reellwertige differenzierbare Funktionen einige zusätzliche Betrachtungen. Beispielsweise lässt sich der Begriff der Differenzierbarkeit direkt auf die Anordnung und Vollständigkeit von $\mathbb{R}$ zurückführen. Übungsaufgabe:

Beispiel 9.5. Für eine reellwertige Funktion einer reellen Veränderlichen soll der Begriff der Ableitung direkt auf die Anordnungs- und Vollständigkeitsstruktur der reellen Zahlen zurückgeführt werden. Sei $x_0 \in I$, einem offenen Intervall, und $f : I \rightarrow \mathbb{R}$. Eine reelle Zahl o heisst *Obersteigung* von f in x_0 , falls ein offenes Intervall $I_o \subset I$ existiert mit $x_0 \in I_o$, sodass für alle $x \in I_o$ gilt, dass

$$\begin{aligned} f(x_0) + o \cdot (x - x_0) &> f(x) \quad \text{falls } x > x_0 \\ \text{und} \quad f(x_0) + o \cdot (x - x_0) &< f(x) \quad \text{falls } x < x_0 \end{aligned} \quad (9.13)$$

Analog definiert man den Begriff einer *Untersteigung*. Sei O die Menge der Obersteigungen von f in x_0 , und U die Menge der Untersteigungen. Überlegen Sie sich zunächst (durch kurze Beweise bzw. geeignete Gegenbeispiele), dass

- (i) O nach oben ordnungsabgeschlossen ist, d.h. $o \in O \Rightarrow o' \in O \forall o' > o$, aber u.U. leer sein kann, und
- (ii) $U < O$, d.h. $u < o$ für alle $u \in U, o \in O$ (also insbesondere ist O nach unten beschränkt, falls $U \neq \emptyset$).
- (iii) Zeigen Sie dann: f ist genau dann in x_0 differenzierbar, wenn O und U nicht-leer sind, und $\inf O = \sup U =: s$. (In diesem Fall gilt $s = f'(x_0)$.)
- (iv) Geben Sie ein Beispiel einer Funktion mit Stelle, für welche U und O nicht leer sind, aber $\inf O \neq \sup U$.

Diese Formulierung ist der eigentliche Ursprung des Zusammenhangs zwischen Differentialrechnung und Monotonie-/Extremwertaufgaben.

Definition 9.6. Man sagt, eine Funktion $F : X \rightarrow \mathbb{R}$ auf einem metrischen Raum X hat in einem Punkt $x_0 \in X$ ein lokales Maximum (d.h., x_0 ist eine lokale Maximumsstelle), falls ein $\delta > 0$ existiert so dass $F(x) \leq F(x_0)$ für alle $x \in B_\delta(x_0)$. Ein lokales Minimum ist analog definiert.

Lemma 9.7. Sei I ein offenes Intervall und $f : I \rightarrow \mathbb{R}$ differenzierbar. Hat f in $x_0 \in I$ ein lokales Maximum oder Minimum, dann gilt $f'(x_0) = 0$.

Beweis. Im Falle eines Maximums existiert ein $\delta > 0$ so, dass $f(x) \leq f(x_0)$ für $x \in I \cap (x_0 - \delta, x_0 + \delta)$. Für die $x > x_0$ ist der Differenzenquotient

$$\frac{f(x) - f(x_0)}{x - x_0} \leq 0 \quad (9.14)$$

und im Grenzfall $x \downarrow x_0$ folgt $f'(x_0) \leq 0$. Durch Betrachtung der $x_0 - \delta < x < x_0$ zeigt man $f'(x_0) \geq 0$, zusammen folgt $f'(x_0) = 0$. Der Fall eines Minimums geht analog. $\square$

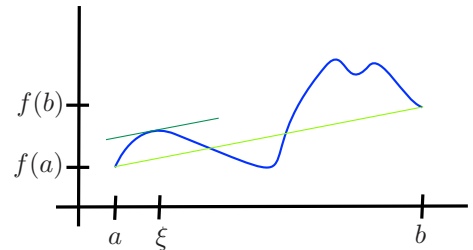
Proposition 9.8. Sei $[a, b] \subset \mathbb{R}$ und $f : [a, b] \rightarrow \mathbb{R}$ stetig, und auf (a, b) differenzierbar. Dann existiert ein $\xi \in (a, b)$ für das $f(b) = f(a) + (b - a)f'(\xi)$.

§9. EINE REELLE VERÄNDERLICHE

Beweis. Wir betrachten die Funktion $g : [a, b] \rightarrow \mathbb{R}$

$$g(t) = f(t) - f(a) - \frac{f(b) - f(a)}{b - a}(t - a) \quad (9.15)$$

g ist auf (a, b) differenzierbar und auf $[a, b]$ stetig mit $g(a) = g(b) = 0$. Nach dem Satz über Maximum und Minimum 7.6 nimmt g ein Maximum M und ein Minimum m an. Gilt $M = m$, so verschwindet g identisch auf $[a, b]$, f ist affin-linear und die Aussage unmittelbar. Andernfalls liegt wegen $g(a) = g(b) = 0$ entweder die Maximumsstelle oder die Minimumsstelle in (a, b) , und wir schliessen mit dem Kriterium 9.7.



□

Definition 9.9. Eine reellwertige Funktion $I \rightarrow \mathbb{R}$ auf einem Intervall heisst konvex, falls für alle $x_1, x_2 \in I$, $x_1 \neq x_2$ und alle $t \in (0, 1)$ gilt

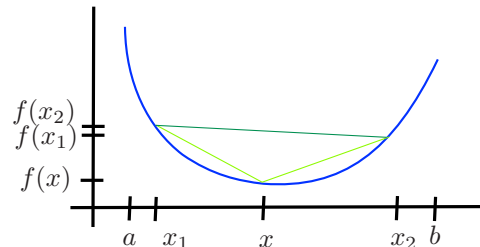
$$f((1-t)x_1 + tx_2) \leq (1-t)f(x_1) + tf(x_2) \quad (9.16)$$

Gilt die strenge Ungleichung, so heisst f streng konvex. Gilt die umgekehrte (strenge) Ungleichung, so heisst f (streng) konkav.

Bemerkungen. Schreibt man $x = (1-t)x_1 + tx_2$, also $t = \frac{x-x_1}{x_2-x_1}$ so wird aus (9.16)

$$\begin{aligned} f(x) &\leq \frac{x_2 - x}{x_2 - x_1} f(x_1) + \frac{x - x_1}{x_2 - x_1} f(x_2) \quad \forall x \in [x_1, x_2] \\ \Leftrightarrow f(x) &\leq f(x_1) + \frac{x - x_1}{x_2 - x_1} (f(x_2) - f(x_1)) \quad \forall x \in [x_1, x_2] \\ \Leftrightarrow f(x) &\leq f(x_2) + \frac{x_2 - x}{x_2 - x_1} (f(x_1) - f(x_2)) \quad \forall x \in [x_1, x_2] \end{aligned} \quad (9.17)$$

“Zwischen je zwei Stellen x_1 und x_2 liegt die Funktion f unterhalb ihrer linearen Interpolation.” Konvexe Funktionen spielen eine sehr wichtige (aber oft versteckte) Rolle, z.B. für die Stabilität physikalischer Systeme. Eine interessante Eigenschaft ist:



Lemma 9.10. Eine konvexe Funktion auf einem offenen Intervall ist stetig.

Beweis. Sei $y_0 \in I$. Da I offen ist, existieren $y_1, y_2 \in I$ mit $y_1 < y_0 < y_2$. Für $y \in [y_0, y_2]$ folgt aus (9.17) (mit $x_1 = y_0, x_2 = y_2, x = y$) einerseits

$$f(y) - f(y_0) \leq \frac{f(y_2) - f(y_0)}{y_2 - y_0} (y - y_0) \quad (9.18)$$

und (mit $x_1 = y_1, x_2 = y, x = y_0$) andererseits

$$f(y) - f(y_0) \geq \frac{f(y_0) - f(y_1)}{y_0 - y_1} (y - y_0) \quad (9.19)$$

Aus dem Sandwich-Lemma folgt $\lim_{y \downarrow y_0} f(y) = f(y_0)$, f ist also rechtsseitig stetig in y_0 . Ebenso zeigt man die linksseitige Stetigkeit, und damit ist f stetig auf ganz I . □

Bei differenzierbaren Funktionen gebraucht man das folgende Ergebnis:

Proposition 9.11. *Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig und auf (a, b) zweimal differenzierbar (d.h. f ist differenzierbar und die Ableitung f' nochmal). Dann ist f genau dann konvex, wenn $f''(x) \geq 0 \ \forall x \in (a, b)$.*

Beweis. Aus den beiden letzten Varianten von (9.17) folgt $\forall x_1, x_2 \in (a, b), x_1 < x_2$:

$$\frac{f(x) - f(x_1)}{x - x_1} \leq \frac{f(x_2) - f(x_1)}{x_2 - x_1} \leq \frac{f(x_2) - f(x)}{x_2 - x} \quad \forall x : x_1 < x < x_2 \quad (9.20)$$

Ist also f konvex, so folgt mit $x \downarrow x_1$ bzw. $x \uparrow x_2$

$$f'(x_1) \leq \frac{f(x_2) - f(x_1)}{x_2 - x_1} \leq f'(x_2) \quad (9.21)$$

d.h. f' wächst monoton auf (a, b) . Daraus folgt durch Grenzwertbildung, dass $f''(x) \geq 0 \ \forall x \in (a, b)$.

- Ist umgekehrt $f'' \geq 0$, so folgt monotonen Wachstum von f' aus dem Mittelwertsatz. (Die Monotonie einer Funktion ist also äquivalent zum Vorzeichen ihrer Ableitung.)
- Für alle $x \in (x_1, x_2)$ folgt aus dem Mittelwertsatz die Existenz von Stellen $\xi_1 \in (x_1, x)$, $\xi_2 \in (x, x_2)$ so, dass

$$f'(\xi_1) = \frac{f(x) - f(x_1)}{x - x_1}, \quad f'(\xi_2) = \frac{f(x_2) - f(x)}{x_2 - x} \quad (9.22)$$

- Aus der Monotonie von f' folgt also

$$\frac{f(x) - f(x_1)}{x - x_1} \leq \frac{f(x_2) - f(x)}{x_2 - x} \quad \forall x_1 < x < x_2 \quad (9.23)$$

was man leicht als äquivalent zur ersten Version von (9.17) erkennt. $\square$

Wir überlassen alle weiteren Feinheiten solcher Kurvendiskussionen den Übungen.¹⁵

Hier irgendwo: Hölder- und Minkowski-Ungleichung

§ 10 Mehrere Veränderliche

Der Ableitungsbegriff aus dem letzten § kontrolliert in linearer (und dann höherer) Näherung die Abhängigkeit (skalarer oder vektorieller) physikalischer Größen von einer einzelnen unabhängigen Variablen, welche am einfachsten als “Zeit” zu interpretieren ist. Wir wollen diesen Begriff nun auf mehrere Veränderliche verallgemeinern, und denken dabei an zwei Typen von physikalischen Beispielen:

- (I) Potentiale $\Phi(x, y, z)$ und Vektorfelder $\vec{F}(x, y, z)$ in Abhängigkeit der drei Ortskoordinaten
- (II) thermodynamische Zustandsfunktionen (Innere Energie, Enthalpie, etc.) eines physikalischen Systems in Abhängigkeit der äusseren Zustandsvariablen (Temperatur, Volumen, etc.)

¹⁵Beispielsweise gilt die Verschärfung von 9.11 nur in der Richtung: $f'' > 0 \Rightarrow f$ streng konvex. f kann streng konvex sein, und f'' trotzdem Nullstellen besitzen. Beispiel: $t \mapsto t^4$ auf $\mathbb{R}$.

§ 10. MEHRERE VERÄNDERLICHE

Im ersten Anlauf können wir solche Abhängigkeiten *von jeder dieser unabhängigen Variablen getrennt* untersuchen, was auf den Begriff der sog. *partiellen Ableitung* führt. Dieser Begriff ist nützlich, aber noch nicht für alle Zwecke befriedigend: In Situationen vom Typ (I) berücksichtigt er nicht die enge Beziehung zwischen den Ortskoordinaten, wie etwa die Möglichkeit der räumlichen Drehungen, die auch noch auf die Komponenten des Vektorfeldes wirken, m.a.W., dass die Identifikation des physikalischen Raums mit $\mathbb{R}^3$ mathematisch einer *Basiswahl* entspricht. Etwas allgemeiner kann man auch an krummlinige Koordinaten denken, relativistisch an die Einbeziehung der Zeit in die Koordinatentransformationen.

Aber auch in Situationen vom Typ (II), in denen die verschiedenen Variablen zunächst eine klar getrennte physikalische Interpretation besitzen (beispielsweise von verschiedener Dimension sind), möchte man manchmal die Rollen von unabhängigen und abhängigen Variablen vertauschen (wie beispielsweise den Druck statt dem Volumen als äussere Variable benutzen) und es erweist sich hierfür als zweckmässig, die Abhängigkeiten von *allen Variablen gemeinsam* zu beschreiben. Dies führt unvermeidlich auf den Begriff der sog. *totalen Ableitung*.

Im Folgenden sind stets $n, m \in \mathbb{N}$, $U \subset \mathbb{R}^n$ offen und $F : U \rightarrow \mathbb{R}^m$ eine (zunächst beliebige) Abbildung.

Definition 10.1. Sei $x_0 = (x_0^i)_{i=1,\dots,n} \in U$. Dann heisst F *partiell differenzierbar* in x_0 falls für jedes $i = 1, 2, \dots, n$ die Einschränkung von F auf die i -te Koordinatenrichtung differenzierbar im Sinne von 9.1 ist, d.h. $\forall i$ existiert die i -te partielle Ableitung

$$\partial_i F(x_0) = \lim_{x^i \rightarrow x_0^i} \frac{F(x_0^1, \dots, x^i, \dots, x_0^n) - F(x_0^1, \dots, x_0^i, \dots, x_0^n)}{x^i - x_0^i} \quad (10.1)$$

- Die Einschränkung auf offene U ermöglicht ähnlich wie in 9.1 die Untersuchung des Grenzwertes “von beiden Richtungen aus”. Etwas allgemeiner könnte man sich aber darauf zurückziehen, dass der Definitionsbereich ein Produkt von evtl. halboffenen Intervallen um x_0 enthält, s. Bemerkungen auf S. 77. (Beispiel: Temperatur T in der Nähe des absoluten Nullpunktes).

- Man beachte, dass wir für die Definition der partiellen Ableitung zwar die Vektorraumstruktur (und Norm) im $\mathbb{R}^m$ benötigen, aber nichts dergleichen im Definitionsbereich. Insbesondere schränkt die Forderung nach der partiellen Differenzierbarkeit nur das Verhalten von F entlang der Koordinatenachsen ein. Man überlegt sich dann auch leicht alle möglichen Beispiele partiell differenzierbarer Funktionen.

- Als erster Schritt zur Befreiung von der Auszeichnung der Koordinatenrichtungen auf $\mathbb{R}^n$ bemerken wir, dass wir den Grenzwert in (10.1) auch schreiben können als

$$\partial_i F(x_0) = \lim_{t \rightarrow 0} \frac{F(x_0^1, \dots, x_0^i + t, \dots, x_0^n) - F(x_0^1, \dots, x_0^i, \dots, x_0^n)}{t} \quad (10.2)$$

Dies (ist informell offensichtlich und) folgt formal z.B. aus dem Folgenkriterium 6.7. (Der Limes über Folgen $x_n^i \rightarrow x_0^i$ ist äquivalent zum Limes über Folgen $t_n = x_n^i - x_0^i \rightarrow 0$.) Wir erklären also

Definition 10.2. Sei $x_0 \in U$, und $v \in \mathbb{R}^n$. Dann existiert wegen der Offenheit von U ein $\delta > 0$ so dass $x_0 + tv \in U$ für alle $t \in (-\delta, \delta)$. F heisst in x_0 *differenzierbar in Richtung v* , falls die Funktion $t \mapsto F(x_0 + tv)$, in $t_0 = 0$ differenzierbar im Sinne von 9.1 ist. Man nennt

$$D_v F(x_0) = \lim_{t \rightarrow 0} \frac{F(x_0 + tv) - F(x_0)}{t} \in \mathbb{R}^m \quad (10.3)$$

die *Richtungsableitung* von F in Richtung v .

- Offensichtlich ist F partiell differenzierbar, wenn sie in die Richtung aller Elemente $e_i = (0, \dots, \overset{i\text{-te Stelle}}{1}, \dots, 0)^T$, $i = 1, 2, \dots, n$ der Standardbasis von $\mathbb{R}^n$ differenzierbar ist. (Es gilt: $\partial_i F(x_0) = D_{e_i} F(x_0)$)
- Beachte, dass man F wegen der Offenheit von U auf Differenzierbarkeit in jede Richtung $v \in \mathbb{R}^n$ untersuchen kann, auch wenn U eine echte Teilmenge des $\mathbb{R}^n$ ist. Geometrisch gesehen ist “der $\mathbb{R}^n$, der U enthält,” der affine Raum im Sinne von 4.1, und der $\mathbb{R}^n$ als “Raum der Richtungen” der zugrundeliegende Vektorraum.
- Wir wollen nun (für festes x_0) die Abhängigkeit von $D_v F(x_0) \in \mathbb{R}^m$ von $v \in \mathbb{R}^n$ untersuchen. Zunächst bemerken wir, dass für alle $\lambda \in \mathbb{R}$,

$$D_{\lambda v} F(x_0) = \lambda D_v F(x_0) \quad (10.4)$$

Bew.: Für $\lambda \neq 0$ folgt dies, auch ohne Kettenregel, aus den Rechenregeln für Folgen, dass nämlich mit $\tilde{t} = \lambda t$

$$\lim_{t \rightarrow 0} \frac{F(x_0 + t\lambda v) - F(x_0)}{t} = \lambda \lim_{\tilde{t} \rightarrow 0} \frac{F(x_0 + \tilde{t}v) - F(x_0)}{\tilde{t}} \quad (10.5)$$

während für $\lambda = 0$ die Aussage trivial ist.

- Die Ableitungen in linear unabhängige Richtungen sind zunächst zwar vollkommen beliebig. Sei beispielsweise $\hat{F} : S^{n-1} \rightarrow \mathbb{R}$ eine beliebige Funktion auf der $(n-1)$ -dimensionalen Einheitssphäre $S^{n-1} = \{\|x\|_2 = 1\} \subset \mathbb{R}^n$. Wählen wir dann für jeden ein-dimensionalen Unterraum (Gerade durch den Ursprung) $\{tv \mid t \in \mathbb{R}\} \neq \{0\} \subset \mathbb{R}^n$ einen der beiden Durchstossunkte durch S^{n-1} als $x_v \in S^{n-1}$ und erklären $F : \mathbb{R}^n \rightarrow \mathbb{R}$ durch

$$F(tx_v) = t\hat{F}(x_v) \quad (10.6)$$

so ist F im Ursprung in jede Richtung differenzierbar mit $D_{x_v} F(0) = \hat{F}(x_v)$.

- Solches beliebiges Verhalten ist aber die Ausnahme. In den meisten Fällen ist die Abbildung $\mathbb{R}^n \ni v \mapsto D_v(x_0) \in \mathbb{R}^m$ *linear in v* , und approximiert die Funktion in der Nähe von x_0 bis auf “Terme höherer als linearer Ordnung”. (Durch geeignete Variationen, z.B. durch Ersetzen $t \rightarrow t^2 \sin \frac{1}{x_v^1 t}$ in (10.6), könnten wir erreichen, dass die Differenzierbarkeit nicht *gleichmässig* im Raum der Richtungen ist, d.h. insbesondere, dass selbst wenn die Abbildung $v \mapsto D_v(x_0)$ linear ist, F noch nicht total differenzierbar sein muss.)

Definition 10.3. Seien $n, m \in \mathbb{N}$ (Dimension 0 ist weniger interessant), und $U \subset \mathbb{R}^n$ offen. Eine Funktion $F : U \rightarrow \mathbb{R}^m$ heisst differenzierbar in $x_0 \in U$, falls eine lineare Abbildung $DF(x_0) : \mathbb{R}^n \rightarrow \mathbb{R}^m$ (das “(totale) Differential von F in x_0 ”) existiert so dass

$$\lim_{x \rightarrow x_0} \frac{\|F(x) - F(x_0) - DF(x_0)(x - x_0)\|}{\|x - x_0\|} = 0 \quad (10.7)$$

§ 10. MEHRERE VERÄNDERLICHE

Bemerkungen. · Hier ist $DF(x_0)(x - x_0)$ die Auswertung der linearen Abbildung $DF(x_0) \in \text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \mathbb{R}^m)$ auf $x - x_0 \in \mathbb{R}^n$. Im Nenner steht eine Norm auf $\mathbb{R}^n$, im Zähler eine auf $\mathbb{R}^m$. Welche Normen zu Grunde liegen spielt wie bereits mehrfach betont keine Rolle, siehe 4.26. (Weder für den Begriff der Differenzierbarkeit noch für den Wert des Differentials, s. unten.)

· Die ein-dimensionale Definition 9.1 benutzt explizit die Körperstruktur von $\mathbb{R}$, um in (9.1) direkt durch den Zuwachs $t - t_0$ zu teilen. Im Unterschied dazu bewertet man in (10.7) die Güte der linearen Approximation über den Vergleich in der Norm. Dennoch reduziert sich natürlich die Definition 10.3 im Fall $n = 1$ wieder auf 9.1. Die Aussage (10.7) ist nämlich äquivalent zu: $\forall \epsilon > 0$ existiert ein $\delta > 0$ so, dass

$$\|F(x) - F(x_0) - DF(x_0)(x - x_0)\| \leq \epsilon \cdot \|x - x_0\| \quad \text{falls } \|x - x_0\| < \delta \quad (10.8)$$

(Für $x \neq x_0$ ist dies klar, für $x = x_0$ wegen des Ersetzens von $<$ durch $\leq$ trivial.)

· Falls F in x_0 differenzierbar ist, so folgt aus der Linearität von $DF(x_0)$ durch Einschränkung des Grenzprozesses $x \rightarrow x_0$ auf die Menge $\{x \in U \mid x - x_0 \in \mathbb{R}v\}$, dass dann F in x_0 auch in alle Richtungen v differenzierbar ist, mit Richtungsableitung $D_v F(x_0) = DF(x_0)(v)$:

$$\begin{aligned} \lim_{t \rightarrow 0} \frac{\|F(x_0 + tv) - F(x_0) - DF(x_0)(tv)\|}{\|tv\|} &= 0 \\ &= \frac{1}{\|v\|} \cdot \left\| \frac{F(x_0 + tv) - F(x_0)}{t} - DF(x_0)(v) \right\| \\ &\Rightarrow \lim_{t \rightarrow 0} \frac{F(x_0 + tv) - F(x_0)}{t} = DF(x_0)(v) \end{aligned} \quad (10.9)$$

Insbesondere ist also das totale Differential durch (10.7) eindeutig bestimmt.

· Genauer gesagt ist $DF(x_0)$ bereits durch die partiellen Ableitungen in alle Koordinatenrichtungen bestimmt, welche gerade die Auswertung $D_{e_i} F(x_0) = DF(x_0)(e_i) \in \mathbb{R}^m$ von $DF(x_0)$ auf der Standard-Basis $(e_1, \dots, e_n)$ von $\mathbb{R}^n$ sind.

· Für $j = 1, \dots, m$ sei F^j die j -te Komponente von F , das heisst für $x \in U$ ist $F(x) = (F^1(x), \dots, F^m(x))^T$. Wie im ein-dimensionalen Fall gilt dann: F ist genau dann differenzierbar, wenn alle Komponenten $F^j : U \rightarrow \mathbb{R}$ differenzierbar sind. (Man verallgemeinert die obigen Begriffe leicht auf unendlich-dimensionale normierte Räume als Wertebereiche.)

Definition 10.4. Die Darstellung der linearen Abbildung $DF(x_0) : \mathbb{R}^n \rightarrow \mathbb{R}^m$ bezüglich den Standard-Basen von $\mathbb{R}^n$ und $\mathbb{R}^m$, d.h. die Matrix der Richtungsableitungen der Komponenten von F heisst Jacobi-Matrix oder Funktionalmatrix von F :

$$DF(x_0)_i^j = \lambda^j(DF(x_0))(e_i) = \frac{\partial F^j}{\partial x^i}(x_0) = \partial_i F^j(x_0) \quad (10.10)$$

($i = 1, \dots, n$; $j = 1, \dots, m$)

Beispiel $F : \mathbb{R}^3 \rightarrow \mathbb{R}^3$, $F(x, y, z) = (x^2y, y^2z, z^2x)^T$:

$$DF(x) = \begin{pmatrix} 2xy & x^2 & 0 \\ 0 & 2yz & y^2 \\ z^2 & 0 & 2zx \end{pmatrix} \quad (10.11)$$

Wir kommen nun zum wichtigsten Kriterium für die totale Differenzierbarkeit.

Definition 10.5. F heisst (partiell) differenzierbar auf U , falls F in jedem $x \in U$ (partiell) differenzierbar ist.

· F heisst stetig partiell differenzierbar auf U , falls F auf U partiell differenzierbar ist und die Abbildungen $x \mapsto \partial_i F(x) \in \mathbb{R}^m$ für $i = 1, \dots, n$ stetig sind.

· F heisst stetig differenzierbar auf U , falls F differenzierbar ist und die Abbildung $U \rightarrow \text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \mathbb{R}^m)$, $x \mapsto DF(x)$ stetig ist.

· In dieser Definition benutzt man auf $\text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \mathbb{R}^m)$ die zu den Normen auf $\mathbb{R}^n$ und $\mathbb{R}^m$ gehörige Operatornorm, vgl. (4.48).

· Nach den obigen Bemerkungen folgt also aus Differenzierbarkeit partielle Differenzierbarkeit. Die Umkehrung gilt aber nur unter der Voraussetzung der Stetigkeit (Gegenbeispiel s. Übungen).

Proposition 10.6. Sei $U \subset \mathbb{R}^n$ offen, und $F : U \rightarrow \mathbb{R}^m$ in einer Umgebung von $x_0 \in U$ stetig partiell differenzierbar. Dann ist F in x_0 differenzierbar.

Beweis. Wegen der letzten Bemerkung vor 10.5 genügt es, den Fall $m = 1$ zu betrachten. Es sei $\delta > 0$ so, dass $\forall x \in B_\delta(x_0)$ der abgeschlossene Quader $\{(1-t)x_0^i + t x^i\}_{i=1,\dots,n} \mid t_i \in [0, 1]\} \subset U$, und alle partiellen Ableitungen $\partial_i F$ dort existieren und stetig sind. Die Funktion $[0, 1] \ni t \mapsto F((1-t)x_0^1 + t x^1, x_0^2, \dots, x_0^n)$ ist dann wohldefiniert und stetig differenzierbar mit Ableitung $\partial_1 F((1-t)x_0^1 + t x^1, x_0^2, \dots, x_0^n)(x^1 - x_0^1)$. (Das sieht man am einfachsten mit der Kettenregel, die wir hier noch nicht haben. Es geht aber auch wie in (10.5).) Daher existiert nach dem (ein-dimensionalen) Mittelwertsatz 9.8 eine Stelle ξ^1 zwischen x_0^1 und x^1 so, dass

$$F(x^1, x_0^2, \dots, x_0^n) - F(x_0^1, x_0^2, \dots, x_0^n) = \partial_1 F(\xi^1, x_0^2, \dots, x_0^n)(x^1 - x_0^1) \quad (10.12)$$

Durch Wiederholung finden wir auch für $i = 2, \dots, n$ Stellen ξ^i zwischen x_0^i und x^i so, dass

$$F(x) - F(x_0) = \sum_{i=1}^n \partial_i F(x^1, \dots, x^{i-1}, \xi^i, x_0^{i+1}, \dots, x_0^n)(x^i - x_0^i) \quad (10.13)$$

Mit der Abkürzung $x_i := (x^1, \dots, x^{i-1}, \xi^i, x_0^{i+1}, \dots, x_0^n)$ gilt dann

$$\begin{aligned} |F(x) - F(x_0) - \sum_{i=1}^n \partial_i F(x_0)(x^i - x_0^i)| &= \left| \sum_{i=1}^n (\partial_i F(x_i) - \partial_i F(x_0))(x^i - x_0^i) \right| \\ &\leq \sum_{i=1}^n |\partial_i F(x_i) - \partial_i F(x_0)| \cdot \|x - x_0\|_\infty \end{aligned} \quad (10.14)$$

Aus der Stetigkeit der $\partial_i F$, der Äquivalenz der Normen, und $\|x_i - x_0\|_\infty \leq \|x - x_0\|_\infty$ folgt, dass die rechte Seite für jedes $\epsilon > 0$ kleiner als $\epsilon \|x - x_0\|$ wird für $\|x - x_0\| < \delta$ klein genug, und damit die Behauptung. $\square$

Beispiel 10.7. Ist $A \in \text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \mathbb{R}^m)$ und $y_0 \in \mathbb{R}^m$, so ist die affin-lineare Abbildung $F(x) := y_0 + A(x)$ differenzierbar mit $DF(x_0) = A \forall x_0 \in \mathbb{R}^n$:

$$\lim_{x \rightarrow x_0} \frac{\|F(x) - F(x_0) - A(x - x_0)\|}{\|x - x_0\|} = \lim_{x \rightarrow x_0} 0 = 0 \quad (10.15)$$

§ 10. MEHRERE VERÄNDERLICHE

- Durch Vergleich von (10.8) mit (9.7) sieht man, dass für $n = 1$ die (gewöhnliche) Ableitung $\dot{f}(t_0)$ (als Vektor in $\mathbb{R}^m$) gleich ist der Auswertung des Differentials $Df(t_0)$ (als lineare Abbildung $\mathbb{R} \rightarrow \mathbb{R}^m$) auf dem Standardbasisvektor von $\mathbb{R}$.
- Es sei $q = (q_{ij})_{i,j=1,\dots,n}$ eine symmetrische Matrix, $q_{ij} = q_{ji}$, $q^T = q^T$, Darstellung einer *quadratischen Form*

$$Q : \mathbb{R}^n \rightarrow \mathbb{R}, \quad x \mapsto \sum_{i,j=1}^n q_{ij} x^i x^j \quad (10.16)$$

Dann ist Q differenzierbar mit

$$DQ(x_0)(v) = 2q(x, v) = 2 \sum_{i,j} q_{ij} x^i v^j \quad (10.17)$$

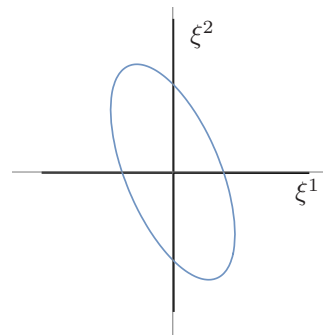
wo $q : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ die zu Q gehörige Bilinearform ist, Beweis über 10.4 oder auch direkt.

· Ein *euklidisches inneres Produkt* auf einem reellen Vektorraum V ist eine positiv definite symmetrische Bilinearform $\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{R}$. Ein *euklidischer Raum* ist ein reeller Vektorraum mit einem euklidischen inneren Produkt. Beispiel: Das Standard-Produkt (4.4) auf $\mathbb{R}^n$.

· Für ein anderes Beispiel betrachten wir auf $\mathbb{R}_\xi^2$ mit Koordinaten (ξ^1, ξ^2)

$$\begin{aligned} \langle \xi_1, \xi_2 \rangle &:= (\xi_1^1, \xi_1^2) \begin{pmatrix} 3 & 1 \\ 1 & 1 \end{pmatrix} \begin{pmatrix} \xi_2^1 \\ \xi_2^2 \end{pmatrix} \\ &= \xi_2^1 (3\xi_1^1 + \xi_1^2) + \xi_2^2 (\xi_1^1 + \xi_1^2) \end{aligned} \quad (10.18)$$

Wegen $\langle \xi, \xi \rangle = 3(\xi^1)^2 + (\xi^2)^2 + 2\xi^1 \xi^2 = 2(\xi^1)^2 + (\xi^1 + \xi^2)^2$ ist dies positiv definit und definiert insbesondere eine Norm auf $\mathbb{R}_\xi^2$. Die zugehörigen “Kugeln” von festem Radius sind Ellipsen, deren genaue Beschreibung wir als Übungsaufgabe lassen.



· Man kann zeigen, dass jeder endlich-dimensionale euklidische Raum isomorph zu $\mathbb{R}^n$ mit dem Standard-Produkt ist. Dies ist gleichbedeutend mit der Existenz einer (nicht eindeutigen) “Orthonormal-”Basis $(e_i)_{i=1,\dots,n}$ mit der Eigenschaft, dass

$$\langle e_i, e_j \rangle = \delta_{ij} = \begin{cases} 1 & \text{für } i = j \\ 0 & \text{sonst} \end{cases} \quad (10.19)$$

Bezüglich einer beliebigen Basis $(\eta_1, \dots, \eta_n)$ des endlich-dimensionalen Vektorraums, mit zugehörigen Koordinaten $(\xi^1, \dots, \xi^n)$ hat das innere Produkt die Form

$$\langle \xi_1, \xi_2 \rangle = \sum_{\alpha, \beta} \xi_1^\alpha g_{\alpha\beta} \xi_2^\beta \quad (10.20)$$

für eine symmetrische positiv definite Matrix $g = (g_{\alpha\beta})$, und der Isomorphismus $B :$

$\mathbb{R}_\xi^n \xrightarrow{\cong} \mathbb{R}^n$ die Darstellung¹⁶

$$\eta_\alpha = \sum_i B_\alpha^i e_i, \quad \text{bzw.} \quad x^i(\xi) = \sum_\alpha B_\alpha^i \xi^\alpha \quad (10.21)$$

mit einer invertierbaren Matrix (B_α^i) , deren Inverses wir als C_i^α schreiben wollen. Es gilt also

$$e_i = \sum_\alpha C_i^\alpha \eta_\alpha, \quad \text{bzw.} \quad \xi^\alpha = \sum_i C_i^\alpha x^i \quad (10.22)$$

In Matrixnotation ist $g = (g_{\alpha\beta}) = B^T \cdot B$.

· Im obigen Beispiel ist $g = \begin{pmatrix} 3 & 1 \\ 1 & 1 \end{pmatrix}$, die Orthonormalbasis $e_1 = \frac{1}{\sqrt{2}}(1, -1)^T$ und $e_2 = (0, 1)^T$, d.h. $C = \begin{pmatrix} \frac{1}{\sqrt{2}} & 0 \\ -\frac{1}{\sqrt{2}} & 1 \end{pmatrix}$ und $B = \begin{pmatrix} \sqrt{2} & 0 \\ 1 & 1 \end{pmatrix}$. Es gilt $C^T \cdot g \cdot C = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$. (Alles Nachprüfen!).

· Allgemein gilt: Für einen endlich-dimensionalen Vektorraum induziert ein euklidisches inneres Produkt einen Isomorphismus zwischen V und dem Dualraum $V^\vee$ von V :

- Für jedes $v \in V$ ist die Abbildung

$$v^\flat : V \rightarrow \mathbb{R}, \quad v^\flat(w) := \langle v, w \rangle \quad (10.23)$$

offensichtlich linear in w , d.h. $v^\flat \in \text{Hom}_{\mathbb{R}}(V, \mathbb{R}) = V^\vee$.

- Ausserdem ist die Zuordnung $v \mapsto v^\flat$ linear in V (per Definition der Rechen-Operationen auf dem Raum der linearen Abbildungen).

- Die Positivität des inneren Produkts impliziert, dass $v \mapsto v^\flat$ injektiv ist: Ist $v^\flat = 0 \in V^\vee$, dann insbesondere auch

$$v^\flat(v) = \langle v, v \rangle = 0 \xrightarrow{\text{pos. def.}} v = 0 \quad (10.24)$$

- Da V endlich-dimensional ist, impliziert der Rangsatz, dass $v \mapsto v^\flat$ ein Isomorphismus von Vektorräumen ist. Wir schreiben das Inverse dieser Abbildung als $V^\vee \ni \lambda \mapsto \lambda^\sharp \in V$.

· Wir erinnern auch daran, dass zu einer Basis $(b_1, \dots, b_n)$ von V die Dualbasis $(\sigma^1, \dots, \sigma^n)$ von $V^\vee$ definiert ist über die Eigenschaft, dass

$$\sigma^i(b_j) = \delta_j^i = \begin{cases} 1 & \text{falls } i = j \\ 0 & \text{sonst} \end{cases} \quad (10.25)$$

· Für den $\mathbb{R}^n$ als Raum der Spaltenvektoren kann der Dualraum $(\mathbb{R}^n)^\vee$ als Raum der Zeilenvektoren aufgefasst werden, die Paarung $(\mathbb{R}^n)^\vee \times \mathbb{R}^n \ni (\lambda, v) \mapsto \lambda(v) \in \mathbb{R}$ als Multiplikation Zeile $\times$ Spalte.

· Der vom Standard-Produkt induzierte Isomorphismus $(\mathbb{R}^n)^\vee \cong \mathbb{R}^n$ ist einfach die Transposition Zeile $\mapsto$ Spalte

$$(\lambda_1, \dots, \lambda_n)^\sharp = (\lambda_1, \dots, \lambda_n)^T = \begin{pmatrix} \lambda_1 \\ \vdots \\ \lambda_n \end{pmatrix} \quad (10.26)$$

· Für allgemeine lineare Koordinaten (vgl. (10.20)) gilt

$$(\lambda^\sharp)^\alpha = \sum_\beta g^{\alpha\beta} \lambda_\beta \quad (10.27)$$

¹⁶Genauer: Die Basis (η_α) stiftet den Isomorphismus $V \cong \mathbb{R}_\xi^n$, die Basis (e_i) den Isomorphismus $V \cong \mathbb{R}^n$. Auf V selbst ist B (per Definition) die Identität.

§ 10. MEHRERE VERÄNDERLICHE

wobei $(g^{\alpha\beta})$ die zu $(g_{\alpha\beta})$ inverse Matrix ist:

$$(g^{\alpha\beta}) = C^T \cdot C \quad (10.28)$$

Beispiel 10.8. $U \subset \mathbb{R}^n$ offen, $f : U \rightarrow \mathbb{R}$. Das (totale) Differential von f ist eine Linearform $Df(x) : \mathbb{R}^n \rightarrow \mathbb{R}$, also ein Element des Dualraums $(\mathbb{R}^n)^\vee$. Man schreibt hier auch häufig df statt Df und für die Dualbasis $(dx^i)_{i=1,\dots,n}$, d.h.

$$df = \sum_i \partial_i f dx^i \quad (10.29)$$

Der Spaltenvektor $\text{grad } f = \begin{pmatrix} \partial_1 f \\ \vdots \\ \partial_n f \end{pmatrix}$ ist gemäss (10.26) der zu df bezüglich dem

Standard-Produkt metrisch äquivalente Vektor in $\mathbb{R}^n \supset U$ (dem Raum der Richtungen in jedem Punkt $x \in U$). Er heisst der *Gradient* von f . Es gilt $\forall v \in \mathbb{R}^n$:

$$\langle \text{grad } f, v \rangle = \langle df^\sharp, v \rangle = df(v) \quad (10.30)$$

was den Gradienten “basisfrei” definiert: Für beliebige lineare Koordinaten auf $\mathbb{R}_\xi^n$ gilt (10.27)

$$\text{grad } f(x)^\alpha = \sum_\beta g^{\alpha\beta} \partial_\beta f(\xi) \quad (10.31)$$

· Während das Differential (10.29) für jede Richtung v den infinitesimalen Zuwachs in f angibt, so bestimmt der Gradient die “Richtung der grössten Steigung” von f , denn für $df(x) \neq 0$ gilt

$$\begin{aligned} df(x) \left(\frac{\text{grad } f(x)}{\|\text{grad } f(x)\|_2} \right) &= \langle \text{grad } f(x), \text{grad } f(x) \rangle^{1/2} \\ &= \max\{df(x)(v) \mid v \in \mathbb{R}^n, \langle v, v \rangle = 1\} \end{aligned} \quad (10.32)$$

(Wegen der Cauchy-Schwarz-Ungleichung (4.9) ist $df(x)(v) = \langle \text{grad } f(x), v \rangle \leq \|\text{grad } f(x)\|_2 \cdot \|v\|_2$ mit Gleichheit genau für v parallel zu $\text{grad } f(x)$.)

Proposition 10.9. Sei $U \subset \mathbb{R}^n$ offen und $F : U \rightarrow \mathbb{R}^m$ in $x_0 \in U$ differenzierbar. Sei sodann $V \subset \mathbb{R}^m$ offen mit $F(x_0) \in V$ und $G : V \rightarrow \mathbb{R}^l$ in $F(x_0)$ differenzierbar. Dann ist $G \circ F$ in x_0 differenzierbar mit

$$D(G \circ F)(x_0) = DG(F(x_0)) \circ DF(x_0) : \mathbb{R}^n \rightarrow \mathbb{R}^l \quad (10.33)$$

Beweis. Die Behauptung macht nur Sinn, wenn $G \circ F$ in einer offenen Umgebung von x_0 definiert ist. Dies folgt aus der Stetigkeit von F in x_0 . ($F^{-1}(V)$ ist offen.)

· Für $\zeta > 0$ gibt es wegen der Differenzierbarkeit von G in $y_0 := F(x_0)$ ein $\epsilon > 0$ so, dass $B_\epsilon(y_0) \subset V$ und

$$\|G(y) - G(y_0) - DG(y_0)(y - y_0)\|_{\mathbb{R}^l} \leq \zeta \cdot \|y - y_0\|_{\mathbb{R}^m} \quad \text{für } y \in B_\epsilon(y_0). \quad (10.34)$$

Wegen der Offenheit von U , der Stetigkeit und der Differenzierbarkeit von F in x_0 gibt es dann ein $\delta > 0$ so, dass gleichzeitig $B_\delta(x_0) \subset U$, $F(B_\delta(x_0)) \subset B_\epsilon(y_0)$ und

$$\|F(x) - F(x_0) - DF(x_0)(x - x_0)\|_{\mathbb{R}^m} \leq \zeta \cdot \|x - x_0\|_{\mathbb{R}^n} \quad \text{für } x \in B_\delta(x_0). \quad (10.35)$$

Dann ist $G \circ F$ für solche x definiert. Wir schätzen ab

$$\begin{aligned} \|G \circ F(x) - G \circ F(x_0) - DG(y_0) \circ DF(x_0)(x - x_0)\|_{\mathbb{R}^l} \\ \leq \|G(F(x)) - G(y_0) - DG(y_0)(F(x) - y_0)\|_{\mathbb{R}^l} \\ + \|DG(y_0)(F(x) - F(x_0) - DF(x_0)(x - x_0))\|_{\mathbb{R}^l} \end{aligned} \quad (10.36)$$

Der erste Term auf der rechten Seite ist wegen $F(x) \in B_\epsilon(y_0)$, (10.34) und (10.35)

$$\leq \zeta \cdot \|F(x) - F(x_0)\|_{\mathbb{R}^m} \leq \zeta \cdot (\|DF(x_0)\| + \zeta) \cdot \|x - x_0\|_{\mathbb{R}^n} \quad (10.37)$$

während der zweite Term wegen (10.35) durch

$$\|DG(y_0)\| \cdot \zeta \cdot \|x - x_0\|_{\mathbb{R}^n} \quad (10.38)$$

abgeschätzt werden kann. (Hier haben wir zweimal die Abschätzung (4.49) (i) von Vektor- durch Operatornorm benutzt.) Zusammen folgt

$$\begin{aligned} \|G \circ F(x) - G \circ F(x_0) - DG(y_0) \circ DF(x_0)(x - x_0)\|_{\mathbb{R}^l} \\ \leq \zeta \cdot (\|DF(x_0)\| + \zeta + \|DG(y_0)\|) \cdot \|x - x_0\|_{\mathbb{R}^n} \end{aligned} \quad (10.39)$$

und damit die Behauptung. $\square$

• Die Formel (10.9) ist ein Beispiel für die Kettenregel: $t \mapsto F(x_0 + tv)$ ist die Verkettung von $t \mapsto x_0 + tv$ mit F und

$$\frac{d}{dt}F(x_0 + tv) = DF(x_0) \circ v = DF(x_0)(v) \quad (10.40)$$

Korollar 10.10. *Ist $U \subset \mathbb{R}^n$ offen und $F : U \rightarrow \mathbb{R}^m$ differenzierbar derart, dass $V := F(U)$ offen ist und eine differenzierbare Abbildung $G : V \rightarrow \mathbb{R}^n$ existiert so dass $G \circ F = \text{id}_U$ und $F \circ G = \text{id}_V$ so gilt $\forall x_0 \in U$ mit $y_0 = F(x_0)$:*

$$DG(y_0) = DF(x_0)^{-1} \quad (10.41)$$

Insbesondere gilt in einem solchen Fall $m = n$. Eine solche Abbildung $F : U \rightarrow V$ heisst Diffeomorphismus.

Lemma 10.11. *Sei $U \subset \mathbb{R}^n$ offen und $F : U \rightarrow \mathbb{R}^m$ stetig differenzierbar. Für je zwei Punkte $x_1, x_2 \in U$, für welche die Verbindungsstrecke $\{(1-t)x_1 + tx_2 \mid t \in [0, 1]\}$ ganz in U liegt, gilt*

$$\|F(x_2) - F(x_1)\| \leq \sup\{\|DF((1-t)x_1 + tx_2)\| \mid t \in [0, 1]\} \cdot \|x_2 - x_1\| \quad (10.42)$$

Insbesondere gilt: Ist $K \subset U$ konvex und folgenkompakt, so ist

$$\|F(x_2) - F(x_1)\| \leq \|DF\|_K \cdot \|x_2 - x_1\| \quad (10.43)$$

Beweis. Die Funktion $f : [0, 1] \rightarrow \mathbb{R}^m$, $f(t) = F((1-t)x_1 + tx_2)$ ist wegen der Kettenregel 10.9 differenzierbar mit $\dot{f}(t) = DF((1-t)x_1 + tx_2)(x_2 - x_1)$ und erfüllt auch die sonstigen Voraussetzungen des Schrankensatzes 9.3. Mit (4.49), d.h. $\|DF(x)(x_2 - x_1)\| \leq \|DF(x)\| \cdot \|x_2 - x_1\|$ folgt die Behauptung. $\square$

Wertebereich $\mathbb{R}$. Höhere Ableitungen

Mit Hilfe von 10.6 und (10.9) lässt sich die Diskussion reellwertiger Funktionen mehrerer Veränderlicher auf die Ableitung nach einer einzelnen Veränderlichen zurückführen.

Lemma 10.12. *Ist $F : U \rightarrow \mathbb{R}$ differenzierbar und x_0 ein lokales Maximum oder Minimum, dann ist $DF(x_0) = 0$.*

Beweis. Es genügt zu zeigen, dass $DF(x_0)(v) = 0 \forall v \in \mathbb{R}^n$. Für $\delta > 0$ klein genug ist $\{x_0 + tv \mid t \in (-\delta, \delta)\} \subset U$ und die Einschränkung $t \mapsto F(x_0 + tv)$ hat ein lokales Maximum (bzw. Minimum) in $t_0 = 0$. Dann folgt die Aussage aus 9.7 und (10.40). $\square$

Wegen der fehlenden Anordnung im Definitionsbereich macht der Begriff der monotonen Funktion für $n > 1$ reelle Variablen keinen Sinn. Demgegenüber ist aber die Diskussion der Konvexität sinnvoll und weiterhin sehr wichtig. Wir zeigen zunächst

Lemma 10.13. *Sei $U \subset \mathbb{R}^n$ offen und $F : U \rightarrow \mathbb{R}$. Angenommen, für ein Paar $i, j \in \{1, \dots, n\}$ existieren $\partial_i F$, $\partial_j F$ und $\partial_{ji} F := \partial_j \partial_i F$ auf U . Ausserdem sei $\partial_{ji} F$ in x_0 stetig. Dann existiert auch $\partial_{ij} F(x_0) = \partial_i \partial_j F(x_0)$ und ist gleich $\partial_{ji} F(x_0)$.*

Beweis. Zur Vereinfachung der Notation sei $n = 2$, $i = 1$, $j = 2$.

Wir zeigen zunächst, welche Sorte Grenzwert der gemischten zweiten Ableitung zu Grunde liegt: Sei für (x^1, x^2) in einem geeigneten Quadrat in U um (x_0^1, x_0^2) herum

$$\Xi(x^1, x^2) := \frac{F(x^1, x^2) - F(x_0^1, x^2) - F(x^1, x_0^2) + F(x_0^1, x_0^2)}{(x^1 - x_0^1)(x^2 - x_0^2)} \quad (10.44)$$

Dann gilt wegen partieller Differenzierbarkeit in die $(1, 0)$ -Richtung

$$\lim_{x^1 \rightarrow x_0^1} \Xi(x^1, x^2) = \frac{\partial_1 F(x_0^1, x^2) - \partial_1 F(x_0^1, x_0^2)}{x^2 - x_0^2} \quad (10.45)$$

und die Voraussetzung über die Existenz der zweiten Ableitung bedeutet

$$\lim_{x^2 \rightarrow x_0^2} \left(\lim_{x^1 \rightarrow x_0^1} \Xi(x^1, x^2) \right) = \partial_{21} F(x_0^1, x_0^2) \quad (10.46)$$

Wir zeigen nun (und eigentlich unabhängig von dieser Betrachtung), dass die umgekehrte Grenzwertreihenfolge das gleiche Ergebnis liefert. Für festes x^2 ist der Zähler von (10.44) als Funktion des ersten Arguments differenzierbar. Wegen des Mittelwertsatzes 9.8 existiert also für festes x^1 eine Stelle ξ^1 zwischen x_0^1 und x^1 derart, dass

$$\Xi(x^1, x^2) = \frac{\partial_1 F(\xi^1, x^2) - \partial_1 F(\xi^1, x_0^2)}{x^2 - x_0^2} \quad (10.47)$$

Da nun $\partial_1 F$ auf U partiell in die $(0, 1)$ -Richtung differenzierbar ist, existiert, wieder nach dem Mittelwertsatz, ein ξ^2 zwischen x_0^2 und x^2 so, dass

$$\Xi(x^1, x^2) = \partial_2(\partial_1 F(\xi^1, \xi^2)) = \partial_{21} F(\xi^1, \xi^2) \quad (10.48)$$

Eine solche Stelle (ξ^1, ξ^2) existiert nun für jedes (x^1, x^2) in einem Quadrat um (x_0^1, x_0^2) , offenbar mit $\|(\xi^1, \xi^2) - (x_0^1, x_0^2)\| \leq C \cdot \|(x^1, x^2) - (x_0^1, x_0^2)\|$ für ein festes $C > 0$. Für gegebenes $\epsilon > 0$ existiert dann wegen der Stetigkeit von $\partial_{21}F$ in (x_0^1, x_0^2) ein $\delta > 0$ so dass

$$|\Xi(x^1, x^2) - \partial_{21}F(x_0^1, x_0^2)| = |\partial_{21}F(\xi^1, \xi^2) - \partial_{21}F(x_0^1, x_0^2)| < \epsilon \quad (10.49)$$

falls $\|x - x_0\| < \delta$.

· Ist nun x^1 fest mit $\|(x^1, 0) - (x_0^1, 0)\| < \delta$, so folgt aus (10.49) im Limes $x^2 \rightarrow x_0^2$ wegen der partiellen Differenzierbarkeit von F in die $(0, 1)$ -Richtung

$$\left| \frac{\partial_2 F(x^1, x_0^2) - \partial_2 F(x_0^1, x_0^2)}{x^1 - x_0^1} - \partial_{21}F(x_0^1, x_0^2) \right| \leq \epsilon \quad (10.50)$$

(Für etwas mehr Transparenz wähle man eine Folge (x_k^2) mit $\|(x^1, x_k^2) - (x_0^1, x_0^2)\| < \delta$ und $x_k^2 \rightarrow x_0^2$, und setze in (10.49) ein.)

· Aus der Aussage (10.50) folgt Existenz und Wert von $\partial_1 \partial_2 F(x_0^1, x_0^2)$. $\square$

Definition 10.14. Sei $k \geq 0$. Eine Funktion $F : U \rightarrow \mathbb{R}$ heisst k mal stetig partiell differenzierbar, wenn alle beliebig gemischten partiellen Ableitungen k -ter Ordnung existieren und stetig sind. Notation für den Vektorraum solcher Funktionen: $\mathcal{C}^k(U)$. Für höher-dimensionalen Wertebereich gilt 10.13 sinngemäss weiter und man schreibt $\mathcal{C}^k(U, \mathbb{R}^m)$.

· Das Differential einer Abbildung in $\mathcal{C}^k(U, \mathbb{R}^m)$ ist eine Abbildung $U \rightarrow \text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \mathbb{R}^m)$ und für $k > 1$ wieder stetig partiell differenzierbar. Das zweite totale Differential hat dann Werte in $\text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \mathbb{R}^m))$, usw., was aber für höhere k immer unübersichtlicher wird.

· Für $m = 1$ aber ist's noch OK, die zweite totale Ableitung hat Werte in $\text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \mathbb{R})) \cong \text{Hom}_{\mathbb{R}}(\mathbb{R}^n \otimes \mathbb{R}^n, \mathbb{R})$ (s. HöMa III für das Tensorprodukt), m.a.W. ist dies eine bilineare Abbildung $\mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$, welche wegen 10.13 symmetrisch ist.

Definition 10.15. Für $F \in \mathcal{C}^2(U) = \mathcal{C}^2(U, \mathbb{R})$ heisst die darstellende Matrix der zweiten totalen Ableitung bezüglich der Standard-Basis im $\mathbb{R}^n$ die "Hesse-Matrix" von F : $H(x) = (H_{ij}(x))_{i,j=1,\dots,n}$ mit

$$H_{ij}(x) = \partial_i \partial_j F(x) \quad (10.51)$$

Es gilt $H(x)^T = H(x) \forall x \in U$.

Die symmetrische Bilinearform ist also (mit Unterdrückung von x) explizit

$$H : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R} \quad (v, w) \mapsto \sum_{i,j=1}^n H_{ij} v^i w^j = H(v, w) := D(DF(v))(w) \quad (10.52)$$

§ 10. MEHRERE VERÄNDERLICHE

Zur Beachtung: Die Hesse-Matrix (H_{ij}) ist darstellende Matrix einer *symmetrischen Bilinearform* auf $\mathbb{R}^n$, was durch die Stellung beider Indizes “unten” angezeigt wird. · Demgegenüber ist für $F \in \mathcal{C}^1(U, \mathbb{R}^m)$ die Jacobi-Matrix $(DF_i^j)_{\substack{i=1,\dots,n \\ j=1,\dots,m}}$ die Matrix-Darstellung einer linearen Abbildung (dem Differential von F), siehe (10.10)

$$DF : \mathbb{R}^n \rightarrow \mathbb{R}^m$$

$$v \mapsto \left(\sum_{i=1}^n DF_i^j v^i \right)_{j=1,\dots,m} = DF(v) \quad (10.53)$$

Für $m = n$ ist diese Matrix auch quadratisch, i.A. aber nicht symmetrisch.

· Man sollte die beiden Begriffe nicht verwechseln.

Definition 10.16. Eine symmetrische Bilinearform $H : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ (oder eine sie darstellende Matrix) heisst *positiv definit*, falls für alle $v \in \mathbb{R}^n \setminus \{0\}$ gilt, dass $H(v, v) > 0$. Sie heisst *positiv semi-definit*, falls $H(v, v) \geq 0$ für alle v . Sie heisst *negativ definit/semi-definit*, falls die umgekehrten Ungleichungen gelten. H heisst *indefinit*, falls sowohl v existieren, für die $H(v, v) > 0$, als auch solche, für die $H(v, v) < 0$.

Sylvester-Kriterium: Eine symmetrische Matrix ist positiv semi-definit genau dann, wenn alle Hauptminoren nicht-negativ sind.

· Sie ist positiv definit genau dann, wenn alle führenden Hauptminoren positiv sind.
· Minoren einer Matrix sind die Determinanten aller quadratischen Untermatrizen, die durch Streichung beliebiger Zeilen und Spalten entstehen. Hauptminoren einer symmetrischen Matrix: Streichung der gleich-indizierten Zeilen und Spalten. Führende Hauptminoren: Determinanten der oberen $k \times k$ Untermatrizen.

Definition 10.17. Eine Teilmenge $U \subset \mathbb{R}^n$ heisst *konvex*, falls $\forall x_1, x_2 \in U$ und alle $t \in (0, 1)$, $(1 - t)x_1 + tx_2 \in U$. Eine Funktion $F : U \rightarrow \mathbb{R}$ auf einer konvexen Menge heisst *konvex*, falls für alle solche x_1, x_2, t gilt, dass $F((1 - t)x_1 + tx_2) \leq (1 - t)F(x_1) + tF(x_2)$.

· Diese Definition stimmt mit der in 9.9 überein, denn für $n = 1$ gilt: Die konvexen Teilmengen von $\mathbb{R}$ sind genau die Intervalle. In Verallgemeinerung von 9.11 gilt:

Proposition 10.18. Sei $U \subset \mathbb{R}^n$ offen und konvex. Dann ist $F \in \mathcal{C}^2(U, \mathbb{R})$ genau dann konvex auf U falls $H(x)$ an jeder Stelle $x \in U$ positiv semi-definit ist.

Beweis. Siehe Lehrbücher. □

Proposition/Definition 10.19. Ist $F : U \rightarrow \mathbb{R}$ differenzierbar, so heisst eine Stelle $x_0 \in U$ mit $DF(x_0) = 0$ ein *kritischer Punkt* von F . Ist $F \in \mathcal{C}^2(U)$, so gilt: Ein kritischer Punkt ist ein *isoliertes lokales Maximum/Minimum*, falls die Hesse-Matrix dort negativ/positiv definit ist.¹⁷ Ist die Hesse-Matrix indefinit, dann ist x_0 kein lokales Extremum.¹⁸

Beweisidee. Benutze die Hesse-Matrix zur quadratischen Approximation der Funktion im Sinne einer mehr-dimensionalen Taylor-Formel. □

¹⁷Die Umkehrung gilt nicht: $x \mapsto x^4$ hat in 0 ein Minimum, die zweite Ableitung verschwindet.

¹⁸D.h., es existieren in jeder Umgebung von x_0 sowohl Punkte x , für die $F(x) > F(x_0)$ als auch solche, für die $F(x) < F(x_0)$.

Divergenz und Rotation

Wie eingangs erwähnt ist, ausser der technischen Spitzfindigkeit, dass nicht jede partiell differenzierbare Funktion total differenzierbar ist, eine wesentliche Motivation für die Einführung des totalen Differentials die gleichberechtigte Behandlung aller unabhängigen Variablen. In der theoretischen Physik kommt diese Sichtweise spätestens in der Hydro- und Elektrodynamik in der Form der Vektordifferentialoperatoren im $\mathbb{R}^3$ zum Tragen. Wir erläutern hier kurz den Zusammenhang zwischen diesen verschiedenen Ableitungsbegriffen, *ohne allerdings auf alle Feinheiten einzugehen*.

Bei der geometrischen Interpretation des Gradienten in 10.8 haben wir darauf hingewiesen, dass seine Definition

$$\text{grad } f := (df)^\# \quad (10.54)$$

von der Auszeichnung eines euklidischen Produkts im “Raum der Richtungen” abhängt. Dieses innere Produkt ist entweder eine eigenständige dynamische Grösse, wie in der Allgemeinen Relativitätstheorie, oder entsteht durch Rückzug des Standardprodukts im $\mathbb{R}^n$ über allenfalls krummlinige Koordinaten. *Die folgenden Formeln gelten aber nur für den Fall, dass die darstellende Matrix $g = (g_{\alpha\beta})$ konstant ist.*

- Der wichtige Punkt der Definition (10.54) ist, dass $\text{grad } f(x)$ für jedes $x \in U$ wieder ein Vektor *im gleichen Vektorraum* ist, in dem auch der ursprüngliche Definitionsbereich lag, nämlich dem Raum der Richtungen an $x \in U \subset \mathbb{R}^n$. Als solcher heisst der $\mathbb{R}^n$ auch “Tangentialraum” an $x \in U$.

- Allgemein heisst eine solche Abbildung $Y : U \rightarrow \mathbb{R}^n$ mit Werten in dem U umgebenden Vektorraum ein (*Tangential-*)*Vektorfeld* auf U .

- Das Differential eines differenzierbaren Vektorfeldes $Y \in \mathcal{C}^1(U, \mathbb{R}^n)$ ist an jeder Stelle $x \in U$ eine lineare Abbildung $DY(x) : \mathbb{R}^n \rightarrow \mathbb{R}^n$. Die *Divergenz* von Y ist die Spur dieser linearen Abbildung. Durch partielle Ableitungen der Komponenten ausgedrückt ist dies einfach

$$\text{div } Y(x) = \text{tr } DY(x) = \sum_{i=1}^n \partial_i Y^i(x) \quad (10.55)$$

- Anschaulich gibt die Divergenz (im geeigneten Kontext) die “Quellendichte” des Vektorfeldes an, was man aber erst durch die Integralsätze wirklich begründen kann. Jedenfalls klingt hier ein Begriff des Volumens an.

- Die Formel (10.55) ist formal unabhängig von g und gilt in jedem linearen Koordinatensystem auf $\mathbb{R}^n$. Ist beispielsweise $x^i(\xi) = \sum_\alpha B_\alpha^i \xi^\alpha$ wie auf Seite 86, so sind die Komponenten von Y bezüglich der Basis (η_α) gegeben durch $Y^\alpha = \sum_i C_i^\alpha Y^i$. Für die Ableitungen gilt gemäss der Kettenregel:

$$\partial_\alpha = \frac{\partial}{\partial \xi^\alpha} = \frac{\partial x^i}{\partial \xi^\alpha} \frac{\partial}{\partial x^i} = B_\alpha^i \partial_i \quad (10.56)$$

Mit $B_i^\alpha C_\beta^i = \delta_\beta^\alpha$ folgt: $\sum_\alpha \partial_\alpha Y^\alpha = \sum_i \partial_i Y^i$ *Beachte:* Für krummlinige Koordinaten muss hier noch die Ableitung des Basiswechsels berücksichtigt werden.

§ 11. KOMPLEX DIFFERENZIERBARE FUNKTIONEN

· Eine spezielle Bildung im drei-dimensionalen Raum ist die *Rotation* eines Vektorfeldes. Diese ist über zwei Ecken definiert als das “Duale des antisymmetrischen Teils des Differentials bzgl. dem orientierten Volumenelement”. Dazu erinnern wir daran, dass das Kreuzprodukt von zwei linear unabhängigen Vektoren $v_1, v_2 \in \mathbb{R}^3$ definiert ist als “derjenige Vektor $v_1 \times v_2 =: v_3 \in \mathbb{R}^3$ mit $\langle v_3, v_1 \rangle = \langle v_3, v_2 \rangle = 0$, $\|v_3\| =$ Flächeninhalt des von v_1, v_2 aufgespannten Parallelogramms, und (v_1, v_2, v_3) positiv orientiert”. Dadurch ist $v_1 \times v_2$ eindeutig bestimmt, ebenso wie $\text{rot } Y$ durch

$$\langle \text{rot } Y(x), v_1 \times v_2 \rangle = \langle DY(x)(v_1), v_2 \rangle - \langle DY(x)(v_2), v_1 \rangle \quad (10.57)$$

für alle $v_1, v_2 \in \mathbb{R}^3$. Für $g = \text{id}$ und Standard-Orientierung ergibt dies die bekannte Formel

$$\text{rot } Y(x) = \begin{pmatrix} \partial_2 Y^3 - \partial_3 Y^2 \\ \partial_3 Y^1 - \partial_1 Y^3 \\ \partial_1 Y^2 - \partial_2 Y^1 \end{pmatrix} \quad (10.58)$$

Anschaulich gibt die Rotation die “gerichtete Wirbeldichte” des Vektorfeldes an. (Die Formel (10.57) berechnet die Wirbeldichte von Y in der von v_1, v_2 aufgespannten Ebene.)

· Im Übrigen gelten die in der theoretischen Physik aufgestellten Rechenregeln und Gesetze. $\text{rot grad } F = 0$, $\text{div rot } Y = 0$. Der Laplace-Operator im $\mathbb{R}^n$ ist

$$\Delta(F) := \text{div grad } F = \sum_{i=1}^n \partial_i^2 F = \sum_{\alpha, \beta=1}^n g^{\alpha\beta} \partial_\alpha \partial_\beta F \quad (10.59)$$

in euklidischen bzw. beliebigen linearen Koordinaten auf $\mathbb{R}^n$.

Fazit: Die in der Physik benutzten Vektordifferentialoperatoren und deren Eigenschaften ergeben sich durch Nachschalten von linearer Algebra im 3-d euklidischen Raum auf die hier vorgestellten mathematischen Begriffe des “totalen Differentials”.

§ 11 Komplex differenzierbare Funktionen

Dieser § wurde aus Zeitgründen übersprungen und im Wintersemester (Höhere Mathematik III) im Wintersemester nachgeholt, s. § 11.

KAPITEL 5

UMKEHRUNG

Wir hatten die Einführung der reellen Zahlen und der Grenzprozesse motiviert durch die Feststellung, dass gewisse natürlich auftretende mathematische Probleme sich mit endlich vielen Rechenschritten höchstens “beliebig genau”, aber nie “exakt” lösen lassen. Mit dem Fundamentalsatz der Algebra hatten wir dann gesehen, dass es in diesem idealisierten Rahmen manchmal möglich ist, die Existenz von Lösungen nachzuweisen, auch ohne ein explizites Verfahren anzugeben, mit dem man solche Lösungen, und sei es nur approximativ, finden kann.

Wir stellen in diesem Kapitel zwei Klassen von Problemen vor, die im Rahmen unserer Differentialrechnung formuliert und lösbar sind. Die Resultate sind zwar einerseits recht allgemeine und abstrakte Existenzaussagen, wie sie auch in der weiteren theoretischen Entwicklung häufig verwendet werden. Die Beweise beruhen aber andererseits auf iterativen Näherungsverfahren, wie sie auch in konkreten Beispielen (zumindestens sinngemäss) implementiert werden können.

Die grundlegende Strategie hinter diesen Lösungsverfahren ist dabei recht simpel. Man identifiziert und löst zunächst eine linearisierte Version des Problems im Rahmen geeignet normierter Vektorräume. Dann stellt man durch geschickte Abschätzungen sicher, dass, ausgehend von einer approximativen Lösung des gegebenen allgemeinen (nicht-linearen) Problems, die Lösung des linearen Problems die approximative Lösung hinreichend verbessert. Zuletzt garantiert dann die Vollständigkeit des umgebenden Raumes, dass die Iteration des linearen Lösungsverfahrens zu einer Lösung des nicht-linearen Problems konvergiert.

§ 12 Determinante und Spur

Zur weiteren Vorbereitung “wiederholen” wir einige Ergebnisse aus der linearen Algebra, insbesondere ein Kriterium für die Lösbarkeit inhomogener linearer Gleichungssysteme der Form

$$(Ax)^j = y^j, \quad j = 1, \dots, n \quad (12.1)$$

für gegebene Matrix $A = (A_i^j) \in \text{Mat}(n, \mathbb{R})$, Spaltenvektor $(y^j) \in \mathbb{R}^n$, und gesuchten Spaltenvektor $(x^i) \in \mathbb{R}^n$.

Lemma/Definition 12.1. *Es existiert eine eindeutige Abbildung $\det : \text{Mat}(n, \mathbb{R}) \rightarrow \mathbb{R}$ (die Determinante) mit den Eigenschaften:*

- (i) *$\det$ ist multi-linear in den Spalten (und Zeilen);*
- (ii) *$\det$ ist alternierend;*
- (iii) *$\det 1 = 1$ (Normierung)*

Beweis. Zur Existenz setzt man

$$\det A = \sum_{\sigma \in S_n} \operatorname{sgn}(\sigma) \prod_{i=1}^n A_{\sigma(i)}^i \quad (12.2)$$

(wo S_n = Permutationsgruppe) und zeigt leicht die Eigenschaften (i), (ii), und (iii). Die Eindeutigkeit führt man durch elementare Zeilen- und Spaltentransformationen auf die Normierung (iii) zurück. $\square$

- Ausser mit der expliziten Form (12.2) lässt sich $\det$ auch durch Entwicklung nach Zeilen und Spalten berechnen.
- Die Determinante ist multiplikativ, das heisst

$$\det(A_1 \cdot A_2) = \det A_1 \cdot \det A_2, \quad \forall A_1, A_2 \in \operatorname{Mat}(n, \mathbb{R}) \quad (12.3)$$

Proposition 12.2. *Sei $A \in \operatorname{Mat}(n, \mathbb{R})$. Dann besitzt die Gleichung $Ax = y$ genau dann für alle $y \in \mathbb{R}^n$ eine eindeutige Lösung $x \in \mathbb{R}^n$, wenn $\det A \neq 0$.*

Beweis. Man zeigt: $\det A \neq 0 \Leftrightarrow A$ ist invertierbar mit inverser Matrix

$$A^{-1} = \frac{1}{\det A} \tilde{A} \quad (12.4)$$

wo sich $\tilde{A}$, die “Adjunkte” zu A , durch die Determinanten

$$\tilde{A}_j^i = (-1)^{i+j} \det A_{\substack{i \\ j}}^{\substack{j \\ i}} \quad (12.5)$$

der Untermatrizen berechnet, die durch Streichung der j -ten Zeile und i -ten Spalte entstehen. (Beachte die Transposition!) $\square$

Beispiel: $n = 2$,

$$A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}, \quad A^{-1} = \frac{1}{ad - bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix} \quad (\text{falls } \det A = ad - bc \neq 0) \quad (12.6)$$

- Die Multiplikativität ist die zentrale Eigenschaft zur Definition der Determinante einer linearen Abbildung.

Lemma/Definition 12.3. *Sei V ein endlich-dimensionaler Vektorraum und $A : V \rightarrow V$ ein Endomorphismus von V . Dann ist die Determinante einer darstellenden Matrix von A bezüglich einer beliebigen Basis von V unabhängig von der Vorgabe dieser Basis und definiert die Determinante von A .*

Beweis. Ist (e_i) eine Basis, aufgefasst als Isomorphismus $V \cong \mathbb{R}_x^n$, so sei (A_j^i) die darstellende Matrix von A . Gibt $B : \mathbb{R}_\xi^n \rightarrow \mathbb{R}_x^n$ einen Basiswechsel auf $\eta_\alpha = \sum_i B_\alpha^i e_i$ (vgl. S. 86), und ist (A_β^α) die darstellende Matrix von A bezüglich (η_α) , so gilt (Lineare Algebra)

$$(A_\beta^\alpha) = B^{-1}(A_j^i)B \quad (12.7)$$

Es folgt:

$$\det(A_\beta^\alpha) = \det B^{-1} \cdot \det(A_j^i) \cdot \det B = \det(A_j^i) \cdot \det(B^{-1}B) = \det(A_j^i) \quad (12.8)$$

$\square$

§ 12. DETERMINANTE UND SPUR

· Im Unterschied dazu ist die “Determinante einer quadratischen Form/eines euklidischen inneren Produkts” $q : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ wesentlich von der Wahl der Basis abhängig. Man setzt wohl *bezüglich der Standardbasis*

$$\det q := \det(q_{ij}) \quad (12.9)$$

aber unter Basiswechsel gilt

$$(q_{\alpha\beta}) = B^T(q_{ij})B \quad (12.10)$$

und daher ist $\det(q_{\alpha\beta})$ nicht gleich $\det(q_{ij})$ im Allgemeinen.

· Aus den gleichen Gründen macht der Begriff der Determinante einer linearen Abbildung zwischen *verschiedenen* Vektorräumen, selbst wenn sie gleiche Dimension haben, *keinen Sinn*. Dennoch gilt auch in diesem Fall

$$\det(\text{darst. Matrix}) \neq 0 \Leftrightarrow \text{Vektorraumisomorphismus} \quad (12.11)$$

· Eine verwandte Idee liegt der Definition der Spur zugrunde: Man erklärt zunächst für eine Matrix $A \in \text{Mat}(n, \mathbb{R})$ durch

$$\text{tr}(A_j^i) := \sum_{i=1}^n A_i^i \in \mathbb{R} \quad (12.12)$$

Dann gilt für $A_1, A_2 \in \text{Mat}(n, \mathbb{R})$ nach der Definition der Matrixmultiplikation:

$$\begin{aligned} \text{tr}(A_1 \cdot A_2) &= \sum_i (A_1 \cdot A_2)_i^i = \sum_i \sum_k A_{1k}^i A_{2i}^k \\ &= \sum_k \sum_i A_{2i}^k A_{1k}^i = \sum_k (A_2 \cdot A_1)_k^k = \text{tr}(A_2 \cdot A_1) \end{aligned} \quad (12.13)$$

und daraus folgt für N Matrizen $A_1, \dots, A_N$ aus der Assoziativität der Matrixmultiplikation,

$$\text{tr}(A_1 \cdot A_2 \cdots A_N) = \text{tr}(A_2 \cdots A_N \cdot A_1) \quad (12.14)$$

Man sagt, die Spur ist “zyklisch symmetrisch”. (Beachte: Im Unterschied zur Determinante darf man unter der Spur nicht beliebig vertauschen.)

· Insbesondere ist

$$\text{tr}(B^{-1}AB) = \text{tr}(ABB^{-1}) = \text{tr} A \quad (12.15)$$

und daher die Definition

$$\text{tr}(\text{Endomorphismus}) = \text{tr}(\text{darstellenden Matrix}) \quad (12.16)$$

unabhängig von der Wahl einer Basis (s. (12.7)).

· Alle diese Definitionen und Eigenschaften gelten auch über einem beliebigen anderen Körper an Stelle von $\mathbb{R}$.

Analytische Eigenschaften. Wir untersuchen nun die Eigenschaften der Determinante als Abbildung des metrischen Raumes/normierten Vektorraums $\text{Mat}(n, \mathbb{R}) \cong \mathbb{R}^{n^2}$ nach $\mathbb{R}$:

- Als erstes halten wir fest, dass die Determinante (12.2) als Polynom in den Standardkoordinaten auf (A_j^i) auf $\mathbb{R}^{n^2}$ eine *stetige* Funktion ist. Insbesondere gilt: Ist $A_0 \in \text{Mat}(n, \mathbb{R})$ invertierbar, was nach 12.2 genau dann zutrifft, wenn $\det A_0 \neq 0$, so existiert für jede Norm $\|\cdot\|$ auf $\text{Mat}(n, \mathbb{R})$ ein $\delta > 0$ s.d. $\det A \neq 0$ falls $\|A - A_0\| < \delta$. Ein solches A ist also auch invertierbar.
- Mit anderen Worten, die Menge der invertierbaren Matrizen,

$$GL(n, \mathbb{R}) = \{A \in \text{Mat}(n, \mathbb{R}) \mid A \text{ invertierbar}\} \quad (12.17)$$

ist eine offene Teilmenge von $\text{Mat}(n, \mathbb{R})$.

- Ist $\|\cdot\|$ speziell eine submultiplikative Norm, so lässt sich die Aussage noch expliziter verifizieren. Beh.: Mit $\delta := \|A_0^{-1}\|^{-1}$ gilt: $A_0 - \Delta$ ist invertierbar falls $\|\Delta\| < \delta$.

Bew.: mit geometrischer Reihe (5.48) in der Banach-Algebra $(\text{Mat}(n, \mathbb{R}), \|\cdot\|)$: Wegen $\|A_0^{-1}\Delta\| \leq \|A_0^{-1}\| \cdot \|\Delta\| < 1$ gilt

$$\sum_{k=0}^{\infty} (A_0^{-1}\Delta)^k = (1 - A_0^{-1}\Delta)^{-1} \quad (12.18)$$

und daher

$$\left(\sum_{k=0}^{\infty} (A_0^{-1}\Delta)^k A_0^{-1}\right) \cdot (A_0 - \Delta) = (1 - A_0^{-1}\Delta)^{-1} (1 - A_0^{-1}\Delta) = 1 \quad (12.19)$$

- Über die Stetigkeit hinaus ist die Determinante als Polynom aber auch (partiell) differenzierbar, und zwar unendlich oft. Wegen 10.6 lässt sich das (totale) Differential also über die Richtungsableitungen ausrechnen. Für $A_0 \in GL(n, \mathbb{R})$ und $\Delta \in \text{Mat}(n, \mathbb{R})$ beliebig ist

$$\lim_{t \rightarrow 0} \frac{1}{t} (\det(A_0 + t\Delta) - \det A_0) = \det A_0 \cdot \lim_{t \rightarrow 0} \frac{1}{t} (\det(1 + t \underbrace{A_0^{-1}\Delta}_{=: X}) - 1) \quad (12.20)$$

und da wie man leicht sieht

$$\det(1 + tX) = \det \begin{pmatrix} 1 + tX_1^1 & \cdots & tX_n^1 \\ \vdots & \ddots & \vdots \\ tX_1^n & \cdots & 1 + tX_n^n \end{pmatrix} = 1 + t \underbrace{\sum_{i=1}^n X_i^i}_{=\text{tr } X} + \mathcal{O}(t^2) \quad (12.21)$$

folgt

$$(D \det)(A_0)(\Delta) = \det A_0 \cdot \text{tr}(A_0^{-1}\Delta) \quad (12.22)$$

“Die Ableitung der Determinante ist die Spur”.

§ 13 Gleichungen im $\mathbb{R}^n$

Bei der ersten Klasse von Problemen ist unsere Aufgabe die Umkehrung einer gegebenen Abbildung $F : U \rightarrow \mathbb{R}^n$ (für $U \subset \mathbb{R}^n$ offen), m.a.W. fragen wir nach der

§ 13. GLEICHUNGEN IM $\mathbb{R}^n$

Lösbarkeit der Gleichung $F(x) = y$ für $y \in \mathbb{R}^n$ und nach den Eigenschaften einer etwaigen Umkehrabbildung in Abhängigkeit derer von F .

Das Kriterium, das wir geben wollen, ist eine mehr-dimensionale Alternative zum strengen Monotonie-Kriterium 7.2 für die Umkehrbarkeit stetiger Funktionen. Dieses Kriterium steht uns ja im Allgemeinen nicht zur Verfügung, da $\mathbb{R}^n$ nicht angeordnet ist. Stattdessen halten wir fest, dass für eine differenzierbare Funktion auf einem Intervall strenge Monotonie gesichert ist, falls die Ableitung ein konstantes Vorzeichen hat, oder als Konsequenz dessen ausgedrückt, als reelle Zahl an jeder Stelle invertierbar ist. Diese Aussage lässt sich nämlich verallgemeinern zur Invertierbarkeit des Differentials von F als lineare Abbildung $DF(x) : \mathbb{R}^n \rightarrow \mathbb{R}^n$.

Ein zweites Zugeständnis müssen wir der Beobachtung aus den Übungen machen, dass für einen allgemeinen Definitionsbereich selbst in dem Fall, in dem eine stetige Funktion global invertierbar ist, die Umkehrfunktion nicht notwendig überall stetig ist. Wir erwarten daher auch im Allgemeinen eine Übertragung der Eigenschaften von F (hier: stetige Differenzierbarkeit) nur nach Einschränkung auf geeignet kleine Umgebungen.

Zur Einstimmung auf die Methode erinnern wir an das Rekursionsverfahren zur Bestimmung von Quadratwurzeln positiver reeller Zahlen (vgl. 3.15). Für $a > 0$ definieren wir eine Abbildung $S^{(a)} : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ durch

$$S^{(a)}(x) := x - \frac{1}{2x}(x^2 - a) = \frac{1}{2}\left(x + \frac{a}{x}\right) \quad (13.1)$$

Mit anderen Worten ist $S^{(a)}$ die Lösung ξ der linearen Gleichung

$$x^2 - a + 2x(\xi - x) = 0 \quad (13.2)$$

oder anschaulich gesagt, die Koordinate des Schnittpunkts der Tangenten an den Graphen von $f^{(a)}(x) := x^2 - a$ im Punkte x mit der x -Achse. Dann gilt $S^{(a)}(x) > \sqrt{a} \forall x > 0$, und $S^{(a)}(x) < x$ falls $x > \sqrt{a}$. Daraus leitet man ab, dass für jeden Startpunkt x_0 die Folge $(x_n := (S^{(a)})^n(x_0))_{n=1,2,\dots}$ monoton fällt und für $n \rightarrow \infty$ gegen $\sqrt{a}$ konvergiert.

Die Konvergenz dieser "Newton-Methode" ist zwar nicht direkt ein Spezialfall des folgenden Beweises, die Stellung der Ableitung $2x = f^{(a)'}(x)$ in (13.1) ist aber analog zu der des Differentials in (13.5) unten. Unser allgemeines Arbeitstier zur Sicherung der Konvergenz ist der Banachsche Fixpunktsatz 4.29, den wir an dieser Stelle kurz wiederholen. Dann aber gehen wir zum

Theorem 13.1. *Sei $U \subset \mathbb{R}^n$ offen und $F : U \rightarrow \mathbb{R}^n$ stetig differenzierbar. Für ein $x_0 \in U$ sei $DF(x_0)$ nicht entartet (d.h., eine invertierbare lineare Abbildung $\mathbb{R}^n \rightarrow \mathbb{R}^n$). Dann existiert eine offene Menge $V_0 \subset \mathbb{R}^n$ mit $F(x_0) \in V_0$, eine offene Menge $U_0 \subset U$ mit $x_0 \in U_0$ und eine stetig differenzierbare Abbildung $G : V_0 \rightarrow U_0$ so, dass $F \circ G = \text{id}_{V_0}$, $G \circ F|_{U_0} = \text{id}_{U_0}$. In Worten: Die Einschränkung von F auf eine hinreichend kleine Umgebung von x_0 ist ein Diffeomorphismus auf das Bild dieser Einschränkung.*

(Für den Begriff des Diffeomorphismusses und die Notwendigkeit der Invertierbarkeit von DF s. 10.10.)

Beweis. Wir konstruieren zunächst eine Umkehrabbildung zu F in einer geeigneten Umgebung von $y_0 := F(x_0)$ mit Hilfe des Banachschen Fixpunktsatzes 4.29 und zeigen anschliessend, dass diese Abbildung differenzierbar ist.

Zur Abkürzung schreiben wir $D_0 := DF(x_0) \in \text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \mathbb{R}^n)$ für das Differential von F in x_0 . Gemäss Annahme existiert $D_0^{-1} : \mathbb{R}^n \rightarrow \mathbb{R}^n$. Wir schreiben die Operatornorm wie auch die (beliebig gewählte) Norm auf $\mathbb{R}^n$, auf der sie beruht, als $\|\cdot\|$.

Auswahl der Umgebungen: Da F in x_0 differenzierbar ist, existiert (zu $\epsilon_* = \frac{1}{2\|D_0^{-1}\|}$) ein $r > 0$ so, dass

$$\|F(x) - F(x_0) - D_0(x - x_0)\| \leq \frac{1}{2\|D_0^{-1}\|} \cdot \|x - x_0\| \quad (13.3)$$

$\forall x \in \overline{B_r(x_0)} \subset U$.¹⁹ Da ferner $DF : U \rightarrow \text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \mathbb{R}^n)$ stetig ist, kann r so gewählt werden, dass ausserdem

$$\|DF(x) - D_0\| < \frac{1}{2\|D_0^{-1}\|} \quad (13.4)$$

$\forall x \in \overline{B_r(x_0)}$. Insbesondere ist $DF(x)$ invertierbar $\forall x \in \overline{B_r(x_0)}$ (geometrische Reihe).

Wir setzen nun $s = \frac{r}{2\|D_0^{-1}\|}$ und $V_0 = B_s(y_0)$.

Konstruktion von G : Für $y \in V_0$ betrachten wir die Abbildung

$$\begin{aligned} T^{(y)} : \overline{B_r(x_0)} &\rightarrow \mathbb{R}^n \\ T^{(y)}(x) &:= x - D_0^{-1}(F(x) - y) \end{aligned} \quad (13.5)$$

(Bemerke, dass Lösungen von $F(x) = y$ genau die Fixpunkte von $T^{(y)}$ sind!) Es gilt

$$\begin{aligned} \|T^{(y)}(x) - x_0\| &= \|x - x_0 - D_0^{-1}(F(x) - F(x_0)) + D_0^{-1}(y_0 - y)\| \\ &\leq \underbrace{\|D_0^{-1}(F(x) - F(x_0) - D_0(x - x_0))\|}_{\leq \|D_0^{-1}\| \cdot \frac{1}{2\|D_0^{-1}\|} \cdot \|x - x_0\|} + \underbrace{\|D_0^{-1}\| \cdot \|y_0 - y\|}_{< \|D_0^{-1}\| \cdot s} \\ &< \frac{r}{2} + \frac{r}{2} = r \end{aligned} \quad (13.6)$$

$T^{(y)}$ bildet also $\overline{B_r(x_0)}$ auf $\overline{B_r(x_0)}$ ab. $T^{(y)}$ ist allerdings auf ganz U differenzierbar mit der Ableitung $\text{id}_{\mathbb{R}^n} - D_0^{-1} \circ DF(x)$, und für $x_1, x_2 \in \overline{B_r(x_0)}$ (einer konvexen Menge) gilt daher wegen des Schrankensatzes 10.11 mit (13.4)

$$\begin{aligned} \|T^{(y)}(x_2) - T^{(y)}(x_1)\| &\leq \sup\{\|\text{id}_{\mathbb{R}^n} - D_0^{-1}DF(x)\| \mid x \in \overline{B_r(x_0)}\} \cdot \|x_2 - x_1\| \\ &\leq \frac{1}{2}\|x_2 - x_1\| \end{aligned} \quad (13.7)$$

$T^{(y)}$ ist also eine Kontraktion auf dem vollständigen metrischen Raum $\overline{B_r(x_0)}$ und hat gemäss 4.29 einen eindeutigen Fixpunkt $G(y)$, der wegen obiger Bemerkung die eindeutige Lösung von $F(x) = y$ in $\overline{B_r(x_0)}$ ist und wegen (13.6) in $B_r(x_0)$ liegt. Dies definiert die Umkehrabbildung $G : V_0 \rightarrow U_0 := F^{-1}(V_0) \cap B_r(x_0)$. ($F^{-1}(V_0)$ ist offen wegen der Stetigkeit von F , und U_0 als Durchschnitt zweier offener Mengen.)

¹⁹Die traditionsgemäss in der offenen Kugel geschriebene Ungleichung (10.8) setzt sich aus Stetigkeitsgründen auf den Abschluss fort, falls dieser in U enthalten ist.

§ 13. GLEICHUNGEN IM $\mathbb{R}^n$

Differenzierbarkeit von G : Es genügt, die Differenzierbarkeit von G in $y_0 = F(x_0)$ nachzuweisen. (Jeder andere Punkt $\tilde{y} \in V_0$ ist ja von der Form $\tilde{y} = F(\tilde{x})$ für ein $\tilde{x} \in B_r(x_0)$ und wegen der Invertierbarkeit von $DF(\tilde{x})$ liefert die obige Konstruktion eine Umkehrung $\tilde{G}$ von F in einer Kugel um $\tilde{y}$. Wegen der Eindeutigkeit der Fixpunkte muss $\tilde{G}$ mit G im Schnitt der beiden Kugeln übereinstimmen. Differenzierbarkeit von $\tilde{G}$ in $\tilde{y}$ impliziert dann die von G .) Wegen der Kettenregel (die wir selbst erst benutzen dürfen, wenn G differenzierbar ist) erwarten wir, dass die Ableitung von G in y_0 durch D_0^{-1} gegeben ist.

Auswerten von (13.7) für $x_1 = x_0$ und $x_2 = x = G(y)$ ergibt

$$\begin{aligned} \|x - x_0\| &= \|T^{(y)}(x) - T^{(y)}(x_0) - D_0^{-1}(y_0 - y)\| \\ &\leq \frac{1}{2}\|x - x_0\| + \|D_0^{-1}\| \cdot \|y - y_0\| \\ \Rightarrow \|x - x_0\| &\leq 2\|D_0^{-1}\| \cdot \|y - y_0\| \end{aligned} \quad (13.8)$$

d.h. $\|G(y) - G(y_0)\| \leq 2\|D_0^{-1}\| \cdot \|y - y_0\| \quad \forall y \in V_0$. Insbesondere ist G stetig in y_0 .

Ist nun $\epsilon > 0$, so existiert wegen der Differenzierbarkeit von F in x_0 ein $\zeta > 0$ so, dass

$$\|F(x) - F(x_0) - D_0(x - x_0)\| \leq \epsilon \cdot \|x - x_0\| \quad \forall x \in B_\zeta(x_0) \subset U_0 \quad (13.9)$$

Diese Aussage ist (trivialerweise) gleichbedeutend mit

$$\|D_0(G(y) - G(y_0)) - (y - y_0)\| \leq \epsilon \cdot \|G(y) - G(y_0)\| \quad \forall y \in G^{-1}(B_\zeta(x_0)) \quad (13.10)$$

Für $\delta \leq \frac{\zeta}{2\|D_0^{-1}\|}$ ist als Konsequenz von (13.8) $G(B_\delta(y_0)) \subset B_\zeta(x_0)$, und es gilt dann $\forall y \in B_\delta(y_0)$:

$$\begin{aligned} \|G(y) - G(y_0) - D_0^{-1}(y - y_0)\| &\leq \|D_0^{-1}\| \cdot \epsilon \cdot \|G(y) - G(y_0)\| \\ &\leq 2\|D_0^{-1}\|^2 \cdot \epsilon \cdot \|y - y_0\| \end{aligned} \quad (13.11)$$

Dies zeigt die Differenzierbarkeit von G in y_0 . Die Stetigkeit von $y \mapsto DG(y) = DF(G(y))^{-1}$ folgt aus der Tatsache, dass dies eine Verkettung der stetigen Funktionen $y \mapsto G(y) = x$, $x \mapsto DF(x)$ und $DF(x) \mapsto (DF(x))^{-1}$ ist. $\square$

Beispiel 13.2. Sei $\Sigma = (\sigma^1, \sigma^2, \sigma^3) : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ definiert durch die *elementarsymmetrischen* Funktionen

$$\begin{aligned} \sigma^1(x^1, x^2, x^3) &= x^1 + x^2 + x^3 \\ \sigma^2(x^1, x^2, x^3) &= x^1x^2 + x^1x^3 + x^2x^3 \\ \sigma^3(x^1, x^2, x^3) &= x^1x^2x^3 \end{aligned} \quad (13.12)$$

(Σ bildet beispielsweise die Wurzeln einer kubischen Gleichung auf ihre Koeffizienten ab. Die Umkehrung dieser Abbildung ist eine natürliche Frage.) Das Differential von Σ bezüglich den Standard-Basen hat die Matrix-Darstellung

$$D\Sigma(x) = \begin{pmatrix} 1 & 1 & 1 \\ x^2 + x^3 & x^1 + x^3 & x^1 + x^2 \\ x^2x^3 & x^1x^3 & x^1x^2 \end{pmatrix} \quad (13.13)$$

mit Determinante

$$\begin{aligned}
 \Delta &= \det D\Sigma(x) = (x^1)^2(x^2 - x^3) - (x^2)^2(x^1 - x^3) + (x^3)^2(x^1 - x^2) \\
 &= (x^1 - x^2)((x^3)^2 + x^1x^2) + (x^2)^2x^3 - (x^1)^2x^3 \\
 &= (x^1 - x^2)((x^3)^2 + x^1x^2 - (x^1 + x^2)(x^3)) \\
 &= (x^1 - x^2)(x^1 - x^3)(x^2 - x^3).
 \end{aligned} \tag{13.14}$$

$D\Sigma(x)$ ist also an allen Stellen $x_0 \in \mathbb{R}^3$ invertierbar, an denen nicht zwei (oder mehr) Komponenten gleich sind. In einer (geeignet kleinen) Umgebung eines jeden solchen Punktes lässt sich Σ also lokal eindeutig und differenzierbar invertieren. Mit anderen Worten: Für gegebene $(\sigma_0^1, \sigma_0^2, \sigma_0^3) \in \Sigma(\mathbb{R}^3)$ gibt es im Allgemeinen zwar mehrere (nämlich 6) mögliche (x^1, x^2, x^3) mit diesen Werten der elementarsymmetrischen Funktionen. Ist aber eine solche Wahl $x_0 = (x_0^1, x_0^2, x_0^3)$ getroffen, so legt diese weitere Wahlen in einer kleinen Umgebung eindeutig fest. Die “schlechten” Werte von $(\sigma^1, \sigma^2, \sigma^3)$ mit weniger als 6 verschiedenen Urbildpunkten sind die Nullstellen von

$$\Delta^2 = (\sigma^1\sigma^2)^2 - 4(\sigma^2)^3 - 4(\sigma^1)^3\sigma^3 + 18\sigma^1\sigma^2\sigma^3 - 27(\sigma^3)^2 \tag{13.15}$$

Global ist Σ weder injektiv noch surjektiv.

· Für eine endlich-dimensionale vollständige normierte Algebra $\mathcal{A}$ (z.B. $\text{Mat}_{n \times n}(\mathbb{R}) \cong \mathbb{R}^{n^2}$ als Vektorraum, versehen mit Matrixmultiplikation und einer submultiplikativen Norm) hatten wir die Exponentialabbildung $\exp : \mathcal{A} \rightarrow \mathcal{A}$ durch die Reihe (5.50) definiert. Wegen $\exp(x)\exp(-x) = 1$ liegt das Bild von $\mathcal{A}$ unter $\exp$ in der Gruppe der invertierbaren Elemente von $\mathcal{A}$. Gemäss (11.18) ist das Differential

$$D\exp(0) = \text{id}_{\mathcal{A}} \tag{13.16}$$

im Ursprung $0 \in \mathcal{A}$ invertierbar. Der Umkehrsatz besagt dann, dass $\exp$ eine geeignete Umgebung von $0 \in \mathcal{A}$ diffeomorph auf eine geeignete Umgebung von $1 \in \mathcal{A}$ abbildet.

Anwendung 13.3 (Krummlinige Koordinaten). Zur Lösung physikalischer Probleme ist das Einführen angepasster Koordinaten oft unumgänglich. Wir beschränken uns hier auf die Vorstellung, dass wir den affinen euklidischen Raum $\mathbb{R}_x^n$ über eine Abbildung

$$K : \mathbb{R}_\xi^n \rightarrow \mathbb{R}_x^n \tag{13.17}$$

durch n Variablen $\xi = (\xi^\alpha) = (\xi^1, \dots, \xi^n)$ parametrisieren wollen. Evtl. schränken wir den Definitionsbereich von K noch ein, setzen aber voraus, dass K mindestens einmal stetig differenzierbar ist.

Beispiel: Lineare Koordinaten wie in (10.21).

Beispiel: Kugelkoordinaten für den $\mathbb{R}^3$:

$$K(r, \theta, \phi) = (r \sin \theta \cos \phi, r \sin \theta \sin \phi, r \cos \theta)^T \tag{13.18}$$

· Wir sagen, Koordinaten (ξ^α) sind “schlecht” an den Bildpunkten $x_0 = K(\xi_0)$, an denen K nicht lokal invertierbar ist, d.h. also falls $\det DK(\xi_0) = 0$.

Im Beispiel:

$$DK = \begin{pmatrix} \sin \theta \cos \phi & r \cos \theta \cos \phi & -r \sin \theta \sin \phi \\ \sin \theta \sin \phi & r \cos \theta \sin \phi & r \sin \theta \cos \phi \\ \cos \theta & -r \sin \theta & 0 \end{pmatrix} \tag{13.19}$$

§ 13. GLEICHUNGEN IM $\mathbb{R}^n$

$$\det DK = r^2 \sin \theta \neq 0 \text{ falls } r \neq 0 \text{ und } \theta \notin \pi\mathbb{Z} \quad (13.20)$$

die schlechten Punkte sind also genau die Menge $\{x = y = 0\}$, wo “ ϕ nicht wohldefiniert ist”.

• Wie oben bemerkt lassen sich solche Koordinate im Allgemeinen nicht global durch eine stetige/differenzierbare Abbildung umkehren. (Im Beispiel liegt dies an $K(r, \theta, \phi) = K(r, \theta, \phi + 2\pi) = K(r, 2\pi - \theta, \phi + \pi)$. Wegen der Lokalität hindert uns dies aber nicht daran, die Differentialrechnung aus dem $\mathbb{R}_x^n$ auf die (ξ^α) zurückzuführen. Wichtiges Beispiel:

Der Laplace-Operator: Wir setzen voraus, dass K zweimal stetig differentierbar ist und arbeiten in offenen Umgebungen, in denen

$$DK = B = \left(B_\alpha^i := \frac{\partial K^i}{\partial \xi^\alpha}(\xi) \right) \quad (13.21)$$

invertierbar ist. Die Notation folgt der aus S. 86, allerdings mit ξ - (oder gleichwertig: x -)abhängigen B und

$$C_i^\alpha = (B^{-1})_i^\alpha, \quad \sum_\alpha B_\alpha^j(K(\xi)) C_i^\alpha(\xi) = \delta_i^j \quad (13.22)$$

Unsere Aufgabe ist das Umschreiben des Ausdrucks

$$\Delta \varphi(x) = \sum_{i=1}^n \frac{\partial^2 \varphi}{\partial x^{i^2}}(x) \quad (13.23)$$

für eine zweimal stetig differenzierbare Funktion $\varphi : \mathbb{R}_x^n \rightarrow \mathbb{R}$ auf die Koordinaten (ξ^α) , d.h. dem Ausdrücken von $\widetilde{\Delta \varphi}(\xi) = \Delta \varphi(K(x))$ durch

$$\tilde{\varphi}(\xi) := \varphi(K(\xi)) \quad (13.24)$$

und ihre partiellen Ableitungen nach ξ .

• Nach der Kettenregel gilt (mit teilweise Unterdrückung der Argumente)

$$\frac{\partial \tilde{\varphi}}{\partial \xi^\alpha} = \sum_{i=1}^n \frac{\partial x^i}{\partial \xi^\alpha} \cdot \frac{\partial \varphi}{\partial x^i} = \sum_i B_\alpha^i \frac{\partial \varphi}{\partial x^i} \quad (13.25)$$

d.h.

$$\frac{\partial \varphi}{\partial x^i} = \sum_\alpha C_i^\alpha \frac{\partial \tilde{\varphi}}{\partial \xi^\alpha} \quad (13.26)$$

und durch nochmalige Anwendung

$$\begin{aligned} \Delta \varphi &= \sum_i \sum_\beta C_i^\beta \frac{\partial}{\partial \xi^\beta} \left(\sum_\alpha C_i^\alpha \frac{\partial \tilde{\varphi}}{\partial \xi^\alpha} \right) \\ (\text{Produktregel}) &= \sum_{i,\alpha,\beta} \frac{\partial}{\partial \xi^\beta} \left(C_i^\beta C_i^\alpha \frac{\partial \tilde{\varphi}}{\partial \xi^\alpha} \right) - \sum_{i,\alpha,\beta} \frac{\partial C_i^\beta}{\partial \xi^\beta} C_i^\alpha \frac{\partial \tilde{\varphi}}{\partial \xi^\alpha} \end{aligned} \quad (13.27)$$

Dieses Umschreiben hat den geometrischen Sinn, dass

$$g^{\beta\alpha} := \sum_i C_i^\beta C_i^\alpha \quad (13.28)$$

als das Inverse der euklidischen Metrik in den ξ -Koordinaten aufgefasst werden kann, vgl. (10.21): $g_{\alpha\beta} = \langle \eta_\alpha, \eta_\beta \rangle$,

$$g = B^T \cdot B, \quad g^{-1} = C \cdot C^T \quad (13.29)$$

(g heisst auch Gramsche Matrix oder erste (Gauss'sche) Fundamentalform.)

· Zur Behandlung des letzten Terms in (13.27) benutzen wir einerseits (12.22) und die Kettenregel:

$$\frac{\partial}{\partial \xi^\beta} \det B = \det B \operatorname{tr} \left(B^{-1} \frac{\partial B}{\partial \xi^\beta} \right) = \det B \sum_\gamma C_j^\gamma \frac{\partial B_j^\gamma}{\partial \xi^\beta} \quad (13.30)$$

und andererseits die Ableitung von (13.22)

$$\sum_\alpha B_\alpha^j \frac{\partial C_i^\alpha}{\partial \xi^\beta} + \sum_\alpha \frac{\partial B_\alpha^j}{\partial \xi^\beta} C_i^\alpha = 0 \quad (13.31)$$

und die Symmetrie der zweiten Ableitung (10.13)

$$\frac{\partial B_\alpha^j}{\partial \xi^\beta} = \frac{\partial^2 K^j}{\partial \xi^\beta \partial \xi^\alpha} = \frac{\partial^2 K^j}{\partial \xi^\alpha \partial \xi^\beta} = \frac{\partial B_\beta^j}{\partial \xi^\alpha} \quad (13.32)$$

um zu schreiben

$$\frac{\partial C_i^\alpha}{\partial \xi^\beta} = - \sum_{j,\gamma} C_j^\gamma \frac{\partial B_\beta^j}{\partial \xi^\gamma} C_i^\alpha \quad (13.33)$$

Zusammen also (unter Umbenennen der Indizes beim Einsetzen von (13.30))

$$\sum_{i,\beta} \frac{\partial C_i^\beta}{\partial \xi^\beta} C_i^\alpha = - \sum_{i,\beta,\gamma} C_j^\beta \frac{\partial B_\beta^j}{\partial \xi^\gamma} C_i^\gamma C_i^\alpha = - \sum_\beta \frac{1}{\det B} \frac{\partial \det B}{\partial \xi^\beta} \cdot g^{\beta\alpha} \quad (13.34)$$

Benutzen wir noch, dass wegen der Multiplikativitat der Determinante und $\det B^T = \det B$, $\det B = \pm \sqrt{\det g}$ (mit einheitlichem Vorzeichen in zusammenhangenden Mengen) so folgt schlussendlich:

$$\begin{aligned} \widetilde{\Delta \varphi} &= \sum_{\alpha,\beta} \left[\frac{\partial}{\partial \xi^\beta} \left(g^{\beta\alpha} \frac{\partial \tilde{\varphi}}{\partial \xi^\alpha} \right) + \frac{1}{\sqrt{\det g}} g^{\beta\alpha} \frac{\partial \sqrt{\det g}}{\partial \xi^\beta} \frac{\partial \tilde{\varphi}}{\partial \xi^\alpha} \right] \\ &= \sum_{\alpha,\beta} \frac{1}{\sqrt{\det g}} \frac{\partial}{\partial \xi^\beta} \left(\sqrt{\det g} g^{\beta\alpha} \frac{\partial \tilde{\varphi}}{\partial \xi^\alpha} \right) \end{aligned} \quad (13.35)$$

In vielen Fallen ist g diagonal, und daher g^{-1} und $\det g$ einfach zu berechnen. Beispiel: Aus (13.19) folgt

$$g = \begin{pmatrix} 1 & 0 & 0 \\ 0 & r^2 & 0 \\ 0 & 0 & r^2 (\sin \theta)^2 \end{pmatrix}, \quad \sqrt{\det g} = r^2 \sin \theta \quad (13.36)$$

und daher ist der Laplace-Operator in 3-d Kugelkoordinaten:

$$\tilde{\Delta} = \frac{1}{r^2} \partial_r r^2 \partial_r + \frac{1}{r^2 \sin \theta} \partial_\theta \sin \theta \partial_\theta + \frac{1}{r^2 \sin^2 \theta} \partial_\phi^2 \quad (13.37)$$

Implizite Funktionen

- Der Satz über die Umkehrfunktion ist eine nicht-lineare Version der Aussage, dass “ n unabhängige Variable durch n unabhängige lineare Gleichungen eindeutig bestimmt sind”, oder in der Sprache der linearen Algebra ausgedrückt, dass für eine nicht-entartete lineare Abbildung $A : \mathbb{R}^n \rightarrow \mathbb{R}^n$ zu jedem $y \in \mathbb{R}^n$ ein eindeutiges x gehört mit $Ax = y$ (nämlich, $x = A^{-1}y$).
- Ebenfalls aus der linearen Algebra bekannt ist die Aussage, dass im Falle von $n - k < n$ unabhängigen Gleichungen, d.h. einer linearen Abbildung $A : \mathbb{R}^n \rightarrow \mathbb{R}^{n-k}$ von maximalem Rang $n - k$, sich die Lösungsmenge der Gleichung $Ax = y$ (für $y \in \mathbb{R}^{n-k}$) parametrisieren lässt durch einen Unterraum von $\mathbb{R}^n$ der Dimension k , dem “Kern” von A . Wir benutzen im Folgenden die Notation

$$\mathbb{R}_{\parallel}^k := \text{Ker } A = \{x_{\parallel} \in \mathbb{R}^n \mid Ax_{\parallel} = 0\} \quad (13.38)$$

und eine Identifikation $\mathbb{R}^n \cong \mathbb{R}_{\parallel}^k \times \mathbb{R}_{\perp}^{n-k}$, welche beispielsweise durch Vervollständigen einer Basis von $\mathbb{R}_{\parallel}^k$ zu einer Basis von $\mathbb{R}^n$ definiert werden kann. Bezüglich dieser Basis hat A die Block-Gestalt

$$A = (0_{(n-k) \times k}, A_{\perp}) \quad (13.39)$$

mit einer invertierbaren $(n - k) \times (n - k)$ Matrix $A_{\perp} \cong A|_{\mathbb{R}_{\perp}^{n-k}}$. Es gilt dann

$$\{x \in \mathbb{R}^n \mid Ax = y\} = \{(x_{\parallel}, (A_{\perp})^{-1}y) \mid x_{\parallel} \in \mathbb{R}_{\parallel}^k\} \quad (13.40)$$

Eine nicht-lineare Version dieser Aussage (mit Unterdrückung von $y = 0$) ist der folgende Satz.

Theorem 13.4. *Seien $n > k$ natürliche Zahlen. Sei $U \subset \mathbb{R}^n$ offen und $F : U \rightarrow \mathbb{R}^{n-k}$ stetig differenzierbar. Sei $x_0 \in U$ eine Lösung der Gleichung $F(x_0) = 0$, in der das Differential $DF(x_0)$ maximalen Rang hat. Wir nehmen oBdA an, dass bezüglich einer Identifikation $\mathbb{R}^n = \mathbb{R}_{\parallel}^k \times \mathbb{R}_{\perp}^{n-k}$, $x = (x_{\parallel}, x_{\perp})$ die Einschränkung $D_{\perp}F(x_0)$ von $DF(x_0)$ auf $\mathbb{R}_{\perp}^{n-k}$ invertierbar ist. Dann existieren offene Umgebungen $V_0 \subset \mathbb{R}_{\parallel}^k$ von $x_{0,\parallel}$ und $W_0 \subset \mathbb{R}_{\perp}^{n-k}$ von $x_{0,\perp}$ mit $V_0 \times W_0 \subset U$, und eine stetig differenzierbare Abbildung $f : V_0 \rightarrow W_0$ so, dass für alle $x = (x_{\parallel}, x_{\perp}) \in V_0 \times W_0$ gilt:*

$$F(x_{\parallel}, x_{\perp}) = 0 \iff x_{\perp} = f(x_{\parallel}) \quad (13.41)$$

M.a.W.,

$$F^{-1}(0) \cap V_0 \times W_0 = \{(x_{\parallel}, f(x_{\parallel})), x \in V_0\} \quad (13.42)$$

In Worten: Die Lösungsmenge von $n - k$ “unabhängigen” differenzierbaren Gleichungen lässt sich lokal als Graph einer Funktion durch k unabhängige Variable parametrisieren.

Beweis. · Die oBdA gemachte Voraussetzung über $D_{\perp}F(x_0)$ ist die Aussage, dass mit der Indizierung $x = (x_{\parallel}, x_{\perp}) = (x_{\parallel}^1, \dots, x_{\parallel}^k, x_{\perp}^{k+1}, \dots, x_{\perp}^n)$ die Untermatrix

$$\left(\frac{\partial F^j}{\partial x_{\perp}^i} (x_0) \right)_{\substack{i=k+1, \dots, n \\ j=1, \dots, n-k}} \quad (13.43)$$

der Jacobi-Matrix von F in x_0 invertierbar ist. Im Allgemeinen kann man dies bereits durch einfaches Umnummerieren der Koordinaten von $\mathbb{R}^n$ erreichen (s. Beispiel unten.)

· Wir erklären nun eine auxiliäre Abbildung $\Phi : U \rightarrow \mathbb{R}_{\parallel}^k \times \mathbb{R}^{n-k}$ durch

$$\Phi(x_{\parallel}, x_{\perp}) := (x_{\parallel}, F(x_{\parallel}, x_{\perp}))^T \quad (13.44)$$

Dann ist Φ auf U stetig differenzierbar, das Differential hat die Block-Gestalt

$$D\Phi(x) = \begin{pmatrix} \text{id}_{k \times k} & 0 \\ D_{\parallel} F(x) & D_{\perp} F(x) \end{pmatrix} \quad (13.45)$$

und ist wegen (13.43) in x_0 invertierbar. Nach dem Umkehrsatz 13.1 existieren daher offene Umgebungen $\tilde{V}_0$ von $\Phi(x_0) = (x_{0,\parallel}, 0)$ und U_0 von x_0 und eine stetig differenzierbare Abbildung $\Psi : \tilde{V}_0 \rightarrow U_0$, welche invers zu $\Phi|_{U_0}$ ist. Aus allgemeinen topologischen Betrachtungen folgt, dass wir oBdA annehmen können, dass $U_0 = U_{0,\parallel} \times U_{0,\perp}$ für offene Umgebungen $U_{0,\parallel}$ von $x_{0,\parallel}$ und $U_{0,\perp}$ von $x_{0,\perp}$. Wir schreiben bezüglich dieser Identifikation $\Psi = (\Psi_{\parallel}, \Psi_{\perp})$.

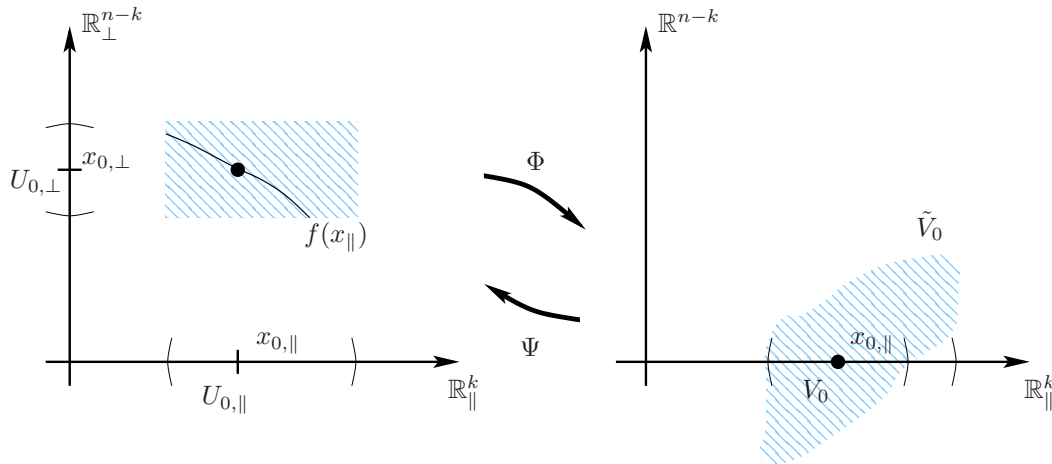
· Wir setzen $V_0 := \{x_{\parallel} \in U_{0,\parallel} \mid (x_{\parallel}, 0) \in \tilde{V}_0\}$ und $W_0 := U_{0,\perp}$. Dann ist V_0 eine offene Umgebung von $x_{0,\parallel}$, und W_0 eine offene Umgebung von $x_{0,\perp}$ und die Definition $f : V_0 \rightarrow W_0$ durch $f(x_{\parallel}) := \Psi_{\perp}(x_{\parallel}, 0)$ macht Sinn.

· f ist stetig partiell differenzierbar da Ψ dies ist, also nach 10.6 stetig differenzierbar.

· Ist für $(x_{\parallel}, x_{\perp}) \in V_0 \times W_0$, $F(x_{\parallel}, x_{\perp}) = 0$, so folgt $(x_{\parallel}, 0) = \Phi(x_{\parallel}, x_{\perp}) \in \tilde{V}_0$, d.h. $(x_{\parallel}, x_{\perp}) = \Psi(x_{\parallel}, 0)$, und daher $x_{\perp} = f(x_{\parallel})$.

· Gilt umgekehrt $x_{\perp} = f(x_{\parallel})$, so ist $(x_{\parallel}, x_{\perp}) = \Psi(x_{\parallel}, 0)$, d.h. $(x_{\parallel}, 0) = \Phi(x_{\parallel}, x_{\perp})$, und daher $F(x_{\parallel}, x_{\perp}) = 0$.

Zusammen folgt (13.41) □



Bemerkungen/Beispiele 13.5. · Aus dem Beweis folgt durch Anwendung der Kettenregel auf die Identität $F(x_{\parallel}, f(x_{\parallel})) = 0$ eine Formel für das Differential der implizit definierten Funktion f :

$$\begin{aligned} D_{\parallel} F(x_{\parallel}, f(x_{\parallel})) + D_{\perp} F(x_{\parallel}, f(x_{\parallel})) \circ Df(x_{\parallel}) &= 0 \\ \Rightarrow Df(x_{\parallel}) &= -D_{\perp} F(x_{\parallel}, f(x_{\parallel}))^{-1} \circ D_{\parallel} F(x_{\parallel}, f(x_{\parallel})) \end{aligned} \quad (13.46)$$

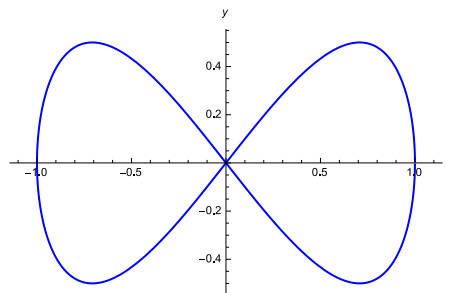
§ 14. STAMMFUNKTIONEN

da $D_\perp F$ wegen der Invertierbarkeit von $D\Phi$ und (13.45) an den angegebenen Stellen invertierbar ist.

· Als Beispiel betrachten wir von der Funktion $F : \mathbb{R}^2 \rightarrow \mathbb{R}$, definiert durch $F(x, y) := x^2(1 - x^2) - y^2$ die Nullstellenmenge $N = \{(x, y) \in \mathbb{R}^2 \mid F(x, y) = 0\}$. Das Differential von F ist dargestellt durch die Jacobi-Matrix

$$DF(x, y) = (2x - 4x^3, -2y) \quad (13.47)$$

Man sieht leicht, dass ausser in $(x, y) = 0 \in N$ DF überall auf N maximalen Rang hat. Allerdings kann man nicht überall mit der gleichen Zerlegung von $\mathbb{R}^2$ in “parallele” und “orthogonale” Richtungen arbeiten. In der Umgebung von $y = 0$ muss man zur Parametrisierung $x_\parallel = y$ benutzen, in der Umgebung von $x = \pm\sqrt{\frac{1}{2}}$ hingegen x .



An allen anderen Punkten kann man N lokal entweder mit der x -Achse oder mit der y -Achse identifizieren. In $(x, y) = 0$ ist überhaupt keine stetige oder differenzierbare ein-dimensionale Parametrisierung von N möglich.

· Der grosse Nutzen des Satzes liegt genau in der soeben illustrierten Möglichkeit, die Parametrisierung der Lösungsmenge zu variieren. Man nennt eine Teilmenge M von $\mathbb{R}^n$, die sich lokal als Nullstellenmenge einer Funktion $F : U \rightarrow \mathbb{R}^{n-k}$ mit DF von maximalem Rang darstellen lässt, eine k -dimensionale “differenzierbare Untermannigfaltigkeit” von $\mathbb{R}^n$. Über eine von 13.4 gelieferte lokale Parametrisierung von M , d.h. $V_0 \xrightarrow{\cong} M \cap V_0 \times W_0$, (die $x_\parallel \mapsto (x_\parallel, f(x_\parallel))$ heissen auch lokale Koordinaten) lässt sich die Differentialrechnung von $V_0 \subset \mathbb{R}_\parallel^k$ auf M hochheben. Dabei wird $\mathbb{R}_\parallel^k$, aufgefasst als Tangentialraum zu V_0 in $x_\parallel$, durch $(\text{id}_{k \times k}, Df(x_\parallel))$ mit einem Unterraum des Tangentialraums zu $\mathbb{R}^n$ in $(x_\parallel, f(x_\parallel))$ identifiziert:

$$T_{x_\parallel} V_0 \cong \mathbb{R}_\parallel^k \ni v \mapsto (v, Df(x_\parallel)v) \in \mathbb{R}^n \cong T_{(x_\parallel, f(x_\parallel))} \mathbb{R}^n \quad (13.48)$$

Im obigen Beispiel ist N (wegen der “Singularität” im Ursprung) keine Untermannigfaltigkeit, $N \setminus \{(0, 0)\}$ hingegen ist.

§ 14 Stammfunktionen

Die zweite Klasse von Problemen, über die wir allgemeine Existenzaussagen machen wollen, sind Anfangswertprobleme für gewöhnliche Differentialgleichungen. Sie haben zahlreiche Beispiele schon in der klassischen Mechanik kennengelernt, wir wiederholen noch einmal die grobe Klassifikation am Anfang des nächsten §.

Zur Vorbereitung führen wir in diesem § das gewöhnliche ein-dimensionale Integral ein zur Lösung der einfachsten Sorte *linearer* Differentialgleichungen. Der entscheidende Schritt zur Lösung gewöhnlicher Differentialgleichungen entlang der allgemeinen Strategie von Seite 95 ist nämlich die Umkehrung der linearen Abbildung, welche einer Funktion $f \in \mathcal{C}^1(I, V)$ ihre Ableitung $\dot{f} \in \mathcal{C}^0(I, V)$ aus §9

zuordnet: Wir hatten schon festgehalten, dass $\mathcal{C}^1(I, V)$ und $\mathcal{C}^0(I, V)$ reelle Vektorräume sind, und die Linearität der Zuordnung $f \mapsto \dot{f}$ ist einfach die Aussage, dass $(f_1 + \lambda f_2) = \dot{f}_1 + \lambda \dot{f}_2 \quad \forall f_1, f_2 \in \mathcal{C}^1(I, V)$ und $\lambda \in \mathbb{R}$.

Für die Abschätzung der “Verbesserung durch Iteration” benötigen wir dann auch geeignete Normen auf $\mathcal{C}^0(I, V)$ und $\mathcal{C}^1(I, V)$. Da diese Vektorräume unendlich-dimensional sind, ist ihre Normierung etwas subtiler als die von $\mathbb{R}^n$, welche wir beim Satz von der Umkehrabbildung benutzt haben. Eigentlich haben wir aber auch schon alle Hilfsmittel im Kapitel 3 vorbereitet. Zunächst einmal benötigen wir, dass für eine stetige Funktion $f : I \rightarrow V$ und jedes kompakte Intervall $[a, b] \subset I$ die Menge $\{\|f(t)\|_V \mid t \in [a, b]\}$ (mit $\|\cdot\|_V$ die Norm auf V) wegen 7.6 beschränkt ist. ($[a, b]$ ist nach dem Satz von Bolzano-Weierstrass 3.14 folgenkompakt und $t \mapsto \|f(t)\|_V \in \mathbb{R}$ ist als Verkettung stetiger Funktionen stetig.) Wir bezeichnen ihr Supremum (in diesem Fall, gleich dem Maximum) mit

$$\|f\|_{[a,b]} := \sup\{\|f(t)\|_V \mid t \in [a, b]\} \quad (14.1)$$

Dann ist die Abbildung $f \mapsto \|f\|_{[a,b]}$ eine Norm auf $\mathcal{C}^0([a, b], V)$.

Definition 14.1. Sei I ein offenes Intervall und $(V, \|\cdot\|)$ ein vollständiger normierter Vektorraum. Eine Funktion $F : I \rightarrow V$ heisst Stammfunktion von $f : I \rightarrow V$, falls F auf I differenzierbar ist und $\dot{F}(t) = f(t) \quad \forall t \in I$.

Bemerkungen. · Eine Stammfunktion ist auf jeden Fall stetig, die gegebene Funktion muss es im Allgemeinen aber nicht sein (vgl. Beispiel (9.5)). Ist jedoch $f \in \mathcal{C}^0(I, V)$, so ist eine Stammfunktion nach dem Schrankensatz 9.3 auf jedem kompakten Teilintervall $K \subset I$ Lipschitz-stetig mit Lipschitz-Konstante $L = \|f\|_K$.

· Sind $F_{1,2}$ Stammfunktionen zu $f_{1,2}$, so ist für $\lambda \in \mathbb{R}$ $F_1 + \lambda F_2$ eine Stammfunktion von $f_1 + \lambda f_2$.

· Gemäss 9.4 ist jede Stammfunktion der Nullfunktion auf einem Intervall konstant. Sind daher F und $\tilde{F}$ zwei Stammfunktionen von f , so gilt $\tilde{F}(t) = F(t) + v_0 \quad \forall t \in I$ für ein $v_0 \in V$.

· Jede jemals errechnete Ableitung gibt auch ein Beispiel einer Stammfunktion. Jedenfalls können Sie schon Polynome, die Exponential- und trigonometrischen Funktion, sowie gewisse Kombinationen dieser Klassen integrieren. Potenzreihen kann man innerhalb ihres Konvergenzintervalls (Schnitt der Konvergenzscheibe mit der reellen Achse) gliedweise integrieren. Im Allgemeinen gilt:

Theorem 14.2. Sei $I \subset \mathbb{R}$ ein offenes Intervall und V ein vollständiger normierter Vektorraum. Dann existiert von jeder stetigen Funktion $f : I \rightarrow V$ eine bis auf die Addition einer Konstanten $v_0 \in V$ eindeutige Stammfunktion.

Die Eindeutigkeit bis auf eine Konstante hatten wir gerade schon festgehalten. Zur Existenz entwickeln wir den Begriff des “bestimmten” Integrals einer stetigen Funktion. Die Intuition dabei ist (*nicht mehr* irgendein “Inhalt der Fläche unter der Kurve”, sondern), dass uns die Ableitung die infinitesimale Veränderung der gesuchten Funktion angibt, wir also nur diese Änderungen über genügend kleine Schritte aufsummieren brauchen.

§14. STAMMFUNKTIONEN

Proposition/Definition 14.3. Sei $I \subset \mathbb{R}$ ein Intervall und $f : I \rightarrow V$ eine stetige Funktion. Für $a, b \in I$ konvergiert die durch

$$S_n(f; a, b) := \sum_{k=0}^{n-1} \left(\frac{b-a}{n} \right) \cdot f\left(a + \frac{k}{n}(b-a)\right) \quad (14.2)$$

definierte Folge $(S_n(f; a, b)) \subset V$. Ihr Grenzwert heisst das Integral von f von a nach b . Schreibweise:

$$\int_a^b f(s) ds := \lim_{n \rightarrow \infty} S_n(f; a, b) \quad (14.3)$$

Für $a = b$ ist offensichtlich $\int_a^a f(s) ds = 0$. Sonst aber erlauben wir in dieser Definition sowohl $a < b$ als auch $b < a$. Aus der Tatsache, dass für alle n

$$\begin{aligned} S_n(f; a, b) &= \sum_{k=0}^{n-1} \left(\frac{b-a}{n} \right) \cdot f\left(\frac{n-k}{n}a + \frac{k}{n}b\right) \\ &= \sum_{k=1}^n \left(-\frac{a-b}{n} \right) \cdot f\left(\frac{k}{n}a + \frac{n-k}{n}b\right) \\ &= -S_n(f; b, a) + \frac{b-a}{n}(f(a) - f(b)) \end{aligned} \quad (14.4)$$

folgt dann sofort eine der wichtigen Rechenregeln für das Integral. Zum Nachweis der Konvergenz von $(S_n(f; a, b))$ holen wir für $a < b$ noch ein wenig weiter aus:

- Eine *Zerlegung* $\mathcal{Z}$ eines Intervalls $[a, b] \subset \mathbb{R}$ ist eine geordnet indizierte endliche Teilmenge von $[a, b]$ mit $a, b \in \mathcal{Z}$, d.h. $\mathcal{Z} = \{a = s_0 < s_1 < \dots < s_r = b\}$ für ein $r \in \mathbb{N}$.

- Die *Feinheit* einer Zerlegung $\mathcal{Z}$ ist die positive Zahl

$$\mu(\mathcal{Z}) := \max\{s_{i+1} - s_i \mid i = 0, \dots, r-1\} \quad (14.5)$$

- Die *Riemannsche Summe*²⁰ von f zur Zerlegung $\mathcal{Z}$ ist

$$S(f; \mathcal{Z}) := \sum_{i=0}^{r-1} (s_{i+1} - s_i) f(s_i) \in V \quad (14.6)$$

Offensichtlich ist (für $a < b$) $S_n(f; a, b)$ aus (14.2) die Riemannsche Summe zur äquidistanten Zerlegung $\{s_k = \frac{n-k}{n}a + \frac{k}{n}b \mid k = 0, \dots, n\}$ der Feinheit $\frac{b-a}{n}$.

- Man sagt, eine Zerlegung $\tilde{\mathcal{Z}}$ ist eine *Verfeinerung* der Zerlegung $\mathcal{Z}$ (Schreibweise: $\tilde{\mathcal{Z}} \triangleleft \mathcal{Z}$), falls $\mathcal{Z} \subset \tilde{\mathcal{Z}}$. Klarerweise gilt in so einem Fall $\mu(\tilde{\mathcal{Z}}) \leq \mu(\mathcal{Z})$.

- Für zwei Zerlegungen $\mathcal{Z}_1$ und $\mathcal{Z}_2$ ist die *gemeinsame Verfeinerung* $\mathcal{Z}_1 \vee \mathcal{Z}_2$ die Vereinigung $\mathcal{Z}_1 \cup \mathcal{Z}_2$ nach einer geeigneten Umindizierung. Es gilt $\mathcal{Z}_1 \vee \mathcal{Z}_2 \triangleleft \mathcal{Z}_1$ und $\mathcal{Z}_1 \vee \mathcal{Z}_2 \triangleleft \mathcal{Z}_2$.

Lemma 14.4. Sei nun $f : [a, b] \rightarrow V$ stetig. Dann existiert für jedes $\epsilon > 0$ ein $\delta > 0$ so, dass für zwei Zerlegungen $\tilde{\mathcal{Z}} \triangleleft \mathcal{Z}$ mit $\mu(\mathcal{Z}) < \delta$

$$\|S(f; \tilde{\mathcal{Z}}) - S(f; \mathcal{Z})\| < \epsilon \quad (14.7)$$

²⁰Dies ist ein Spezialfall der Riemannschen Summen. Im Allgemeinen wertet man f an beliebigen "Stützstellen" $\tau_i \in [s_i, s_{i+1}]$ aus.

Beweis. Wir nutzen aus, dass gemäss 8.9 f als stetige Funktion auf dem folgenkompakten Intervall $[a, b]$ gleichmässig stetig ist. Es existiert also zu gegebenem $\epsilon > 0$ ein $\delta > 0$ so, dass

$$\|f(t_1) - f(t_2)\| < \frac{\epsilon}{b-a} \quad \forall t_1, t_2 \in [a, b] \text{ mit } |t_1 - t_2| < \delta \quad (14.8)$$

Sei nun $\mathcal{Z} = \{a = s_0 < s_1 < \dots < s_r = b\}$ eine Zerlegung von $[a, b]$ mit $\mu(\mathcal{Z}) < \delta$ und $\tilde{\mathcal{Z}} = \{a = \tilde{s}_0 < \tilde{s}_1 < \dots < \tilde{s}_{\tilde{r}} = b\} \triangleleft \mathcal{Z}$ eine Verfeinerung von $\mathcal{Z}$. Dann gehört zu jedem $\tilde{s}_j \in \tilde{\mathcal{Z}}$ ein eindeutiges $s_{i(j)} \in \mathcal{Z}$ mit $s_{i(j)} \leq \tilde{s}_j < s_{i(j)+1}$, und man überlegt sich (am besten mit einer Skizze), dass

$$(i) \quad s_{i+1} - s_i = \sum_{j \mid i(j)=i} (\tilde{s}_{j+1} - \tilde{s}_j) \quad \forall i = 0, \dots, r-1 \quad (14.9)$$

$$(ii) \quad |\tilde{s}_j - s_{i(j)}| < \delta \quad \forall j = 0, 1, \dots, \tilde{r} \quad (14.10)$$

Daraus folgt

$$\begin{aligned} \|S(f; \tilde{\mathcal{Z}}) - S(f; \mathcal{Z})\| &= \left\| \sum_{j=0}^{\tilde{r}-1} (\tilde{s}_{j+1} - \tilde{s}_j) f(\tilde{s}_j) - \sum_{i=0}^{r-1} (s_{i+1} - s_i) f(s_i) \right\| \\ &\quad (\text{mit (i)}) = \left\| \sum_{j=0}^{\tilde{r}-1} (\tilde{s}_{j+1} - \tilde{s}_j) (f(\tilde{s}_j) - f(s_{i(j)})) \right\| \\ &\quad (\text{wegen (ii) und (14.8)}) < \sum_{j=0}^{\tilde{r}-1} (\tilde{s}_{j+1} - \tilde{s}_j) \frac{\epsilon}{b-a} = \epsilon \end{aligned} \quad (14.11)$$

□

Korollar 14.5. Für jedes $\epsilon > 0$ existiert ein $\delta > 0$ so, dass für je zwei Zerlegungen $\mathcal{Z}_1$ und $\mathcal{Z}_2$ mit $\mu(\mathcal{Z}_1) < \delta$ und $\mu(\mathcal{Z}_2) < \delta$

$$\|S(f; \mathcal{Z}_1) - S(f; \mathcal{Z}_2)\| < \epsilon \quad (14.12)$$

Beweis. Sei $\delta > 0$ nach 14.4 so, dass $\|S(f; \tilde{\mathcal{Z}}) - S(f; \mathcal{Z})\| < \frac{\epsilon}{2}$ für $\tilde{\mathcal{Z}} \triangleleft \mathcal{Z}$ und $\mu(\mathcal{Z}) < \delta$. Dann gilt für $\mu(\mathcal{Z}_{1,2}) < \delta$:

$$\begin{aligned} \|S(f; \mathcal{Z}_1) - S(f; \mathcal{Z}_2)\| &\leq \|S(f; \mathcal{Z}_1) - S(f; \mathcal{Z}_1 \vee \mathcal{Z}_2)\| + \|S(f; \mathcal{Z}_1 \vee \mathcal{Z}_2) - S(f; \mathcal{Z}_2)\| < \epsilon \end{aligned} \quad (14.13)$$

□

Beweis von 14.3. Für $a < b$ folgt aus den obigen Betrachtungen, dass $S_n(f; a, b)$ eine Cauchy-Folge ist (s. Def. 4.9): Wähle für $\epsilon > 0$ mit $\delta > 0$ gemäss 14.5 N_ϵ so, dass $\frac{b-a}{N_\epsilon} < \delta$. Dann folgt $\|S_n(f; a, b) - S_m(f; a, b)\| < \epsilon$ für alle $n, m \geq N_\epsilon$. Wegen der Vollständigkeit von V existiert also der Grenzwert $\lim_{n \rightarrow \infty} S_n(f; a, b)$.

• Für $b < a$ folgt die Behauptung dann aus (14.4). □

§ 14. STAMMFUNKTIONEN

Lemma 14.6. (i) Orientierungswechsel: $\int_a^b f(s)ds = - \int_b^a f(s)ds$

(ii) Linearität und Integration von Konstanten: Für $f_1, f_2 \in \mathcal{C}^0(I, V)$ und $\lambda \in \mathbb{R}$, $v \in V$ gilt

$$\int_a^b (f_1 + \lambda f_2 + v)(s)ds = \int_a^b f_1(s)ds + \lambda \int_a^b f_2(s)ds + (b-a)v$$

(iii) Standardabschätzungen ($a < b$): $\left\| \int_a^b f(s)ds \right\| \leq \int_a^b \|f(s)\|ds \leq (b-a) \cdot \|f\|_{[a,b]}$

(iv) Intervalladditivität: Für $a, b, c \in I$ gilt

$$\int_a^b f(s)ds + \int_b^c f(s)ds = \int_a^c f(s)ds$$

Beweis. (i) folgt aus (14.4), (ii) und (iii) aus den entsprechenden Eigenschaften der endlichen Summen (14.2).

(iv) Durch Ausnutzen von (i) können wir die Aussage auf den Fall $a < b < c$ zurückführen. Für $\epsilon > 0$ sei $\delta > 0$ gemäss 14.5 so, dass $\|S(f; \mathcal{Z}_1) - S(f; \mathcal{Z}_2)\| < \epsilon$ für je zwei Zerlegungen von $[a, c]$ mit $\mu(\mathcal{Z}_{1,2}) < \delta$. Sei dann N so, dass $S_n(f; a, b)$, $S_n(f; b, c)$ und $S_n(f; a, c)$ für $n \geq N$ in der Norm weniger als ϵ von den entsprechenden Integralen abweichen, sowie $\frac{c-a}{N} < \delta$. Dann sind $S_n(f; a, b) + S_n(f; b, c)$ und $S_n(f; a, c)$ Riemannsche Summen zu Zerlegungen von $[a, c]$ der Feinheit kleiner als δ . Es folgt für $n \geq N$:

$$\begin{aligned} \left\| \int_a^b f(s)ds + \int_b^c f(s)ds - \int_a^c f(s)ds \right\| &\leq \left\| \int_a^b f(s)ds - S_n(f; a, b) \right\| \\ &+ \left\| \int_b^c f(s)ds - S_n(f; b, c) \right\| + \left\| \int_a^c f(s)ds - S_n(f; a, c) \right\| \\ &+ \|S_n(f; a, b) + S_n(f; b, c) - S_n(f; a, c)\| \\ &< 4\epsilon \end{aligned} \quad (14.14)$$

□

Beweis von 14.2 (Hauptsatz der Differential- und Integralrechnung). Wähle $t_* \in I$ beliebig und setze für $t \in I$

$$F(t) := \int_{t_*}^t f(s)ds \quad (14.15)$$

Beh.: für alle $t_0 \in I$ ist $\dot{F}(t_0) = f(t_0)$.

Bew.: Für $\epsilon > 0$ sei $\delta > 0$ so, dass $\|f(s) - f(t_0)\| < \epsilon$ falls $|s - t_0| < \delta$ (Stetigkeit von f). Dann folgt unter Ausnutzung von 14.6, (ii), (i), (iv), und (iii) für alle $|t - t_0| < \delta$

$$\begin{aligned} \|F(t) - F(t_0) - (t - t_0)f(t_0)\| &= \left\| \int_{t_*}^t f(s)ds - \int_{t_*}^{t_0} f(s)ds - \int_{t_0}^t f(t_0)ds \right\| \\ &= \left\| \int_{t_0}^t (f(s) - f(t_0))ds \right\| \\ &\leq |t - t_0| \cdot \epsilon \end{aligned} \quad (14.16)$$

Die Behauptung folgt in der Formulierung (9.7). □

Korollar 14.7. Sei $f : I \rightarrow V$ stetig und F eine Stammfunktion von f . Dann gilt für alle $a, b \in I$:

$$\int_a^b f(s)ds = F(b) - F(a) \quad (14.17)$$

Beweis. Für beliebiges $t_* \in I$ ist $t \mapsto \int_{t_*}^t f(s)ds$ nach dem Hauptsatz (14.16) eine weitere Stammfunktion von f . Wegen der Eindeutigkeit der Stammfunktion (s.S. 108) gilt dann für alle $t \in I$

$$F(t) = \int_{t_*}^t f(s)ds + v_0 \quad (14.18)$$

für ein festes $v_0 \in V$. Dann folgt für alle $a, b \in I$ wegen 14.6 (i) und (iv):

$$\int_a^b f(s)ds = \int_{t_*}^b f(s)ds - \int_{t_*}^a f(s)ds = F(b) + v_0 - F(a) - v_0 = F(b) - F(a) \quad (14.19)$$

□

· Die äquidistanten Riemannschen Summen (14.2) können zur approximativen Berechnung von Integralen benutzt werden (besser ist die Anpassung der Feinheit wie in (14.6)), die Auswertung des Limes (14.3) gelingt aber nur selten. Üblicherweise “rät” man eine Stammfunktion oder schaut sie in einer vorsorglich angelegten Tabelle nach, oder wendet einen der bekannten Rechentricks an. Die folgenden zwei folgen als Umkehrung von Produkt- bzw. Kettenregel:

Partielle Integration: Ist auf V eine Multiplikation definiert, $F \in \mathcal{C}^1(I, V)$ eine Stammfunktion von $f \in \mathcal{C}^0(I, V)$, und $G \in \mathcal{C}^1(I, V)$, so ist

$$t \mapsto F(t) \cdot G(t) - \int_{t_*}^t F(s) \cdot \dot{G}(s)ds \quad (14.20)$$

eine Stammfunktion von $f \cdot G$.

Substitutionsregel: Ist $F \in \mathcal{C}^1(I, V)$ eine Stammfunktion von $f \in \mathcal{C}^0(I, V)$ und $\tau : J \rightarrow I$ stetig differenzierbar, dann ist $F \circ \tau$ eine Stammfunktion von $(f \circ \tau)\tau'$. Der Umgang dieser Regel mit den Integrationsgrenzen erfordert etwas Konzentration: Für $a, b \in I$ und $\alpha, \beta \in J$ mit $a = \tau(\alpha)$, $b = \tau(\beta)$ gilt:

$$\int_a^b f(s)ds = \int_\alpha^\beta f(\tau(\sigma))\tau'(\sigma)d\sigma \quad (14.21)$$

“Man ersetzt $s = \tau(\sigma)$ und $ds = \frac{ds}{d\sigma}d\sigma = \tau'(\sigma)d\sigma$ und findet Stellen $\alpha, \beta \in J$ so, dass s von a nach b läuft, während σ von α nach β läuft.” Interessant ist, dass in (14.21) τ nicht injektiv sein muss, und noch nicht einmal $\tau([\alpha, \beta]) \subset [a, b]$ gelten muss. Die “überschüssigen” Teile heben sich wegen (14.6) gegenseitig auf. Beispiel: $s = \sigma^2$ in

$$\int_1^2 sds = \frac{1}{2}t^2 \Big|_1^2 = \frac{3}{2} = \int_{-1}^{\sqrt{2}} \sigma^2 \cdot 2\sigma \cdot d\sigma \quad (14.22)$$

§ 14. STAMMFUNKTIONEN

Allerdings ist es nicht möglich, die Anordnung der Zerlegungsstellen in den Riemannschen Summen (14.6) aufzugeben: Betrachte etwa $f(t) = t$ auf dem Intervall $[0, 1]$ und für $n \in \mathbb{N}$ die “Zitterzerlegung” (*kein* sanktionierter Begriff)

$$(s_i) = \left(0, \underbrace{\frac{1}{n}, \frac{2}{n}, \dots, 0, \frac{1}{n}, \frac{2}{n}, \frac{3}{n}, \frac{4}{n}, \dots, 1}_{n^2 \text{ Mal } 3} \right) \quad (14.23)$$

(mit $r = 3n^2 + n - 2$ Elementen). Es gilt $\lim_{n \rightarrow \infty} \max\{|s_{i+1} - s_i| \mid i = 0, \dots, r-1\} = 0$ aber

$$\begin{aligned} \sum_i (s_{i+1} - s_i) f(s_i) &= (n^2 - 1) \cdot \left(\frac{1}{n} \cdot \frac{1}{n} - \frac{2}{n} \cdot \frac{2}{n} \right) + \sum_{k=0}^{n-1} \frac{1}{n} \cdot \frac{k}{n} \\ &= \frac{-3(n^2 - 1)}{n^2} + \frac{1}{n^2} \cdot \frac{n(n-1)}{2} \xrightarrow{n \rightarrow \infty} -3 + \frac{1}{2} \neq \int_0^1 s ds \end{aligned} \quad (14.24)$$

Das Problem mit den Zerlegungen (14.23) ist, dass ihre absolute Länge, nämlich $\sum |s_{i+1} - s_i| = (n^2 - 1) \frac{4}{n} + 1 \rightarrow \infty$ nicht beschränkt bleibt, sie in gewissem Sinne also gar nicht “von a nach b kommen”.

Parameterabhängige Integrale

Ein weiterer häufig benutzter (und auch oft übersehener) Trick bei der Auswertung bestimmter Integrale ohne Auffinden einer Stammfunktion ist die Untersuchung der Abhängigkeit von zusätzlichen Parametern im Integranden. Grundlage dafür ist das folgende Resultat.

Proposition 14.8. *Sei $[a, b] \subset \mathbb{R}$, $U \subset \mathbb{R}^n$ offen und $F : [a, b] \times U \rightarrow V$ stetig. Dann ist die Funktion*

$$G : U \rightarrow V, \quad G(x) := \int_a^b F(s, x) ds \quad (14.25)$$

stetig. Ist V endlich-dimensional, für jedes $t \in [a, b]$ die Abbildung $x \mapsto F(t, x)$ auf U differenzierbar in Richtung $v \in \mathbb{R}^n$, und die Abbildung $[a, b] \times U \ni (t, x) \mapsto D_v F(t, x) \in V$ stetig, dann ist G auf U ebenfalls differenzierbar in Richtung v mit Richtungsableitung

$$D_v G(x) = \int_a^b D_v F(s, x) ds \quad (14.26)$$

Ausserdem übertragen sich diese Aussagen sinngemäss auf (totale) Differenzierbarkeit sowie auf Holomorphie. Die Aussage zur Stetigkeit bleibt gültig, wenn man U durch einen beliebigen metrischen Raum ersetzt.

Beweis. · Die erste Voraussetzung besagt, dass F als Funktion der beiden Variablen gemeinsam stetig ist. Insbesondere ist aber für jedes $x \in U$ die Abbildung $t \mapsto F(t, x)$ stetig, und daher G wohldefiniert, ebenso wie die rechte Seite in (14.26).

· Für $x_0 \in U$ folgt als Korollar von 8.9 aus der Folgenkompaktheit von $[a, b]$, dass für jedes $\epsilon > 0$ ein $\delta > 0$ existiert so, dass $\|F(t, x) - F(t, x_0)\| < \frac{\epsilon}{b-a}$ für alle $t \in [a, b]$

falls nur $\|x - x_0\| < \delta$.²¹ Daraus folgt mit 14.6, dass für $\|x - x_0\| < \delta$

$$\begin{aligned} \|G(x) - G(x_0)\| &= \left\| \int_a^b (F(s, x) - F(s, x_0)) ds \right\| \\ &\leq (b - a) \|F(\cdot, x) - F(\cdot, x_0)\|_{[a, b]} < \epsilon \end{aligned} \quad (14.27)$$

Dies bedeutet, dass G in x_0 stetig ist.

· Mit Hilfe des Mittelwertsatzes 9.8 (getrennt angewendet auf jede der endlich vielen Komponenten von F) und der Stetigkeit von $D_v F$ folgert man recht analog, dass ein $\delta > 0$ existiert so, dass

$$\|F(t, x_0 + \tau v) - F(t, x_0) - D_v F(t, x_0) \tau\| < \frac{\epsilon}{b - a} \cdot |\tau| \quad (14.28)$$

für alle $t \in [a, b]$ falls nur $|\tau| < \delta$. Daraus folgt durch eine Abschätzung wie in (14.27) die Behauptung zur Differenzierbarkeit von G . $\square$

· Als Beispiel behandeln wir das Integral

$$G = \int_0^1 \frac{s^2 - 1}{\ln s} ds \quad (14.29)$$

(Der Integrand ist an den Grenzen zunächst nicht definiert, lässt sich stetig fortsetzen. Siehe auch uneigentliche Integrale weiter unten.) Wir identifizieren $G = G(2)$, wobei

$$G(x) = \int_0^1 \frac{s^x - 1}{\ln s} ds \quad (14.30)$$

und der Integrand für $x \geq 0$ die Voraussetzungen von 14.8 erfüllt. (Hier könnte man $G(x)$ als uneigentliches Integral auf $x > -1$ fortsetzen.) Es gilt dann

$$\begin{aligned} G'(x) &= \int_0^1 \frac{s^x \ln s}{\ln s} ds \\ &= \int_0^1 s^x ds = \frac{s^{x+1}}{x+1} \Big|_0^1 = \frac{1}{x+1} \end{aligned} \quad (14.31)$$

Wegen $G(0) = 0$ folgt daraus durch Integration

$$G(x) = \ln(x + 1) \quad (14.32)$$

und daher $G = G(2) = \ln 3$.

²¹In unserer Herangehensweise reicht es für diese Aussage nicht vorauszusetzen, dass $F(t, \cdot)$ für alle t in x_0 stetig ist, sondern es muss tatsächlich F als Funktion von (t, x) in einer Umgebung von $[a, b] \times \{x_0\}$ stetig sein. Ebenso reicht es für (14.28) nicht aus vorauszusetzen, dass $F(t, \cdot)$ für alle t in x_0 differenzierbar ist.

Uneigentliche Integrale

Ist $I = (\alpha, \beta)$ ein offenes, evtl. auch unbeschränktes Intervall (d.h. $\alpha, \beta \in \mathbb{R} \cup \{\pm\infty\}$), so nennt man eine stetige Funktion $f : (\alpha, \beta) \rightarrow \mathbb{R}$ ("uneigentlich") integrierbar über (α, β) , falls für ein $t_0 \in I$ die rechts-/linksseitigen Limites

$$\lim_{a \downarrow \alpha} \int_a^{t_0} f(s) ds \quad \text{und} \quad \lim_{b \uparrow \beta} \int_{t_0}^b f(s) ds \quad (14.33)$$

beide und getrennt existieren. Man schreibt in diesem Fall weiter

$$\int_{\alpha}^{\beta} f(s) ds = \lim_{a \downarrow \alpha} \int_a^{t_0} f(s) ds + \lim_{b \uparrow \beta} \int_{t_0}^b f(s) ds \quad (14.34)$$

denn das Resultat ist unabhängig von t_0 und stimmt mit der alten Definition überein, falls f in α und/oder β stetig fortsetzbar ist.

- Es existieren Kriterien, die die Existenz solcher uneigentlicher Integrale sicherstellen und gegebenenfalls die Differentiation unter dem Integralzeichen rechtfertigen. (Insbesondere: Majorantenkriterium) Beispiele:
- Die Γ -Funktion ist für $x > 0$ definiert durch

$$\Gamma(x) := \int_0^{\infty} s^{x-1} e^{-s} ds \quad (14.35)$$

und stellt dort eine differenzierbare Funktion dar. Definiert man für $z \in \mathbb{C}$, $s > 0$ $s^z = \exp(z \ln s)$, so gibt für $\operatorname{Re}(z) > 0$

$$\Gamma(z) := \int_0^{\infty} s^{z-1} e^{-s} ds \quad (14.36)$$

eine holomorphe Fortsetzung der Γ -Funktion auf die rechte Halbebene.

- Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion, von der bekannt ist: (i) f hat endlich viele Nullstellen $x_1, \dots, x_r$ mit $f'(x_i) \neq 0$ an jeder dieser Stellen, und (ii) $f(x) \geq O(x)$ für $x \rightarrow \infty$ (d.h. für x gross genug ist $|f(x)| \geq C \cdot |x|$ für eine geeignete Konstante $C > 0$).

Beh.: Dann existiert

$$G = \int_{-\infty}^{\infty} f'(x) e^{-f(x)^2} dx \quad (14.37)$$

und es gilt:

$$G = \sum_i \operatorname{sgn} f'(x_i) \cdot \int_{-\infty}^{\infty} e^{-x^2} dx = \sum_i \operatorname{sgn} f'(x_i) \cdot \sqrt{\pi} \quad (14.38)$$

Bew.: (Skizze) Für $a > 0$ erfüllt der Integrand von

$$G(a) = \int_{-\infty}^{\infty} a f'(x) e^{-a^2 f(x)^2} dx \quad (14.39)$$

die Voraussetzungen zur Differentiation unter dem Integral und es gilt

$$\begin{aligned} G'(a) &= \frac{dG}{da}(a) = \int_{-\infty}^{\infty} (f'(x) - 2a^2 f'(x) f(x)^2) e^{-a^2 f(x)^2} dx \\ &= \int_{-\infty}^{\infty} \frac{d}{dx} (f(x) e^{-a^2 f(x)^2} dx) \\ &= 0 \end{aligned} \quad (14.40)$$

$G(a)$ ist also konstant gleich $G(1) = G$. Andererseits zeigt man unter den gegebenen Voraussetzungen, dass für jedes $\epsilon > 0$

$$\left| G(a) - \sum_{i=1}^r \int_{-\infty}^{\infty} a f'(x_i) e^{-a^2 f'(x_i)^2 (x-x_i)^2} dx \right| < \epsilon \quad (14.41)$$

für a gross genug. (Der Punkt ist, dass der Integrand von $G(a)$ ausserhalb von immer kleiner werdenden Umgebungen der x_i gegen 0 strebt.) Mit der Substitution $x \rightarrow x_i + \frac{x}{a|f'(x_i)|}$ folgt

$$\int_{-\infty}^{\infty} a f'(x_i) e^{-a^2 f'(x_i)^2 (x-x_i)^2} dx = \frac{f'(x_i)}{|f'(x_i)|} \cdot \int_{-\infty}^{\infty} e^{-x^2} dx \quad (14.42)$$

und daraus die Behauptung. Der Beweis von (14.41) wird nachgeliefert. Für die Berechnung des Gaussischen Integrals, siehe §§ 18.

§ 15 Gewöhnliche Differentialgleichungen

Problemstellung: · Gegeben sind:

- $\mathbb{R}^N$ ein endlich-dimensionaler Vektorraum mit Norm $\|\cdot\|$.
- Ein (Zeit-)Intervall $I \subset \mathbb{R}$
- Eine stetige Funktion (Vektorfeld) $Y : I \times U \rightarrow \mathbb{R}^N$ auf dem Produkt von I mit einer offenen Menge $U \subset \mathbb{R}^N$.
- Eine Startzeit $t_0 \in I$ und ein Startwert $x_0 \in U$
- Gesucht sind: ein Intervall $[t_0, t_1] \subset I$ und eine stetig differenzierbare Funktion $x : [t_0, t_1] \rightarrow U$, für die gilt:

$$\dot{x}(t) = Y(t, x(t)) \quad \forall t \in [t_0, t_1] \quad \text{und} \quad x(t_0) = x_0 \quad (15.1)$$

Man nennt (15.1) ein explizites²² Anfangswertproblem (AWP) erster Ordnung, die Funktion x eine Lösung des AWP.

Fragen: · Ist das AWP (15.1) überhaupt lösbar?

- Ist die Lösung eindeutig, zumindest in dem Sinne, dass falls $\tilde{x} : [t_0, \tilde{t}_1] \rightarrow U$ eine zweite Lösung ist, $x(t) = \tilde{x}(t) \quad \forall t \in [t_0, \min(t_1, \tilde{t}_1)]$?
- Wie weit lässt sich die Lösung maximal fortsetzen? Läuft sie etwa vor dem Ende der Zeit aus U heraus?

²²Dies bedeutet, dass die Gleichung nach $\dot{x}(t)$ aufgelöst gegeben ist. Bei impliziten Differentialgleichungen der Form $F(\dot{x}(t), x(t)) = 0$ kann man mal zunächst den Satz über implizite Funktionen 13.4 anrufen.

§ 15. GEWÖHNLICHE DIFFERENTIALGLEICHUNGEN

· Hängt die Lösung in stetiger und/oder differenzierbarer Weise von den Anfangswerten, sowie gegebenenfalls noch von weiteren “externen” Parametern ab? Mit anderen Worten, existiert mit einer offenen Umgebung V von x_0 eine stetige Funktion $\Phi : [t_0, t_1] \times V \rightarrow U$ mit

$$\frac{\partial \Phi}{\partial t}(t, x) = Y(t, \Phi(t, x)), \quad \Phi(t_0, x) = x \quad (15.2)$$

Ist Φ differenzierbar, falls Y dies ist?

Beispiel 15.1. · Siehe theoretische Physik-Vorlesungen!

· Tatsächlich ist Stetigkeit des Vektorfeldes hinreichend für die lokale Existenz von Lösungen. Wir werden diese Aussage nicht beweisen, da Stetigkeit für die Eindeutigkeit im Allgemeinen nicht ausreicht. Als Gegenbeispiel dient hier üblicherweise das Anfangswertproblem mit $I = \mathbb{R}$, $N = 1$, $Y(x) = \sqrt{|x|}$, $t_0 = 0$, $x_0 = 0$, oder kürzer

$$\dot{x} = \sqrt{|x|}, \quad x(0) = 0 \quad (15.3)$$

Offensichtlich ist $x(t) = 0 \forall t$ eine Lösung, aber auch für jedes $t_* \geq 0$:

$$x(t) = \begin{cases} 0 & t \leq t_* \\ \frac{(t - t_*)^2}{4} & \text{für } t > t_* \end{cases} \quad (15.4)$$

Die fehlende Eindeutigkeit hat damit zu tun, dass Y zwar stetig, aber nicht Lipschitzstetig ist, siehe Voraussetzung in 15.2. unten).

· Ein Standardbeispiel für endliche Existenz ist das AWP

$$\dot{x} = x^2, \quad x(0) = x_0 > 0 \quad (15.5)$$

mit Lösung

$$x(t) = \frac{1}{x_0^{-1} - t} \quad t \in [0, x_0^{-1}) \quad (15.6)$$

· Das AWP (15.5) ist vom Typ der “getrennten Veränderlichen”. Allgemein hat (hoffentlich bekanntlich) für zwei offene Intervalle $I, U \subset \mathbb{R}$, eine stetige Funktion $f : I \rightarrow \mathbb{R}$ und eine stetige Funktion $g : U \rightarrow \mathbb{R}$ mit $g(x) \neq 0 \forall x \in U$ das AWP

$$\dot{x}(t) = f(t)g(x(t)), \quad x(t_0) = x_0 \quad (15.7)$$

($t_0 \in I$, $x_0 \in U$) die Lösung

$$x(t) = H^{-1}(F(t)) \quad (15.8)$$

wobei $F = \int_{t_0}^t f(s)ds$ die Stammfunktion von f mit $F(t_0) = 0$ ist und H^{-1} die Umkehrfunktion der Stammfunktion

$$H(x) = \int_{x_0}^x \frac{1}{g(y)} dy \quad (15.9)$$

von $1/g$ mit $H(x_0) = 0$: H ist wegen $\frac{1}{g} \neq 0$ streng monoton und daher nach 7.2 global invertierbar. Die Lösung (15.8) ist definiert auf dem Teilintervall J von

$$\{t \in I \mid F(t) \in H(U)\} = F^{-1}(H(U)) \quad (15.10)$$

welches t_0 enthält. ($F^{-1}(H(U))$ ist offen, da $H(U)$ als monotones Bild eines offenen Intervalls offen ist, und F stetig ist.)

· Ebenfalls bekannt sein sollte die Rückführung von Differentialgleichungen höherer Ordnung auf die Form (15.1). So ist etwa für $F : \mathbb{R}^k \rightarrow \mathbb{R}$ die gewöhnliche skalare Differentialgleichung

$$y^{(k)}(t) = F(y(t), \dot{y}(t), \dots, y^{(k-1)}(t)) \quad (15.11)$$

äquivalent zu einer vektoriellen Gleichung für $x = (y, \dot{y}, \dots, y^{(k-1)})^T \in \mathbb{R}^k$:

$$\frac{d}{dt} \begin{pmatrix} y(t) \\ \dot{y}(t) \\ \vdots \\ y^{(k-1)}(t) \end{pmatrix} = \begin{pmatrix} \dot{y}(t) \\ \ddot{y}(t) \\ \vdots \\ F(y(t), \dots, y^{(k-1)}(t)) \end{pmatrix} \quad (15.12)$$

nämlich, $\dot{x} = Y(x(t))$, wobei $Y(x^1, \dots, x^k) = (x^2, \dots, F(x^1, \dots, x^k))^T$. (Anfangswerte und explizite Zeitabhängigkeit wurden unterdrückt.)

· Eine andere “Reduktion der Ordnung”, welche in der Physik eine grosse Rolle spielt, ergibt sich aus Erhaltungssätzen. Beispielsweise ist auf Lösungen von

$$\ddot{x}(t) = -\text{grad } V(x(t)) \quad (15.13)$$

die Energie $E = \frac{1}{2} \langle \dot{x}, \dot{x} \rangle + V(x)$ konstant. (Wobei grad gemäss (10.54) bezüglich eines euklidischen inneren Produkts $\langle \cdot, \cdot \rangle$ auf $\mathbb{R}^N$ definiert ist, und $V \in C^1(\mathbb{R}^N, \mathbb{R})$, hier also kein Vektorraum.) Im Falle $N = 1$ erhält man dann durch Auflösen nach $\dot{x}$ eine Gleichung mit getrennten Veränderlichen und Lösungen der Form

$$t - t_0 = \int_{x_0}^x \frac{1}{\sqrt{2(E - V(x))}} dx \quad (15.14)$$

(Konvergenz solcher Integrale als Übungsaufgabe.)

· Für eine (fixierte) lineare Abbildung $A : \mathbb{R}^N \rightarrow \mathbb{R}^N$ hat für jedes $x_0 \in \mathbb{R}^N$ das AWP

$$\dot{x}(t) = Ax(t), \quad x(0) = x_0 \quad (15.15)$$

die Lösung

$$x(t) = \exp(tA)x_0 \quad (15.16)$$

Beispiel: Schwingungsgleichung. Zurück zur Theorie.

Theorem 15.2. *Seit $I \subset \mathbb{R}$ ein offenes Intervall, $U \subset \mathbb{R}^N$ offen und $Y : I \times U \rightarrow \mathbb{R}^N$ stetig. Angenommen, es existiert eine Lipschitz-Konstante $L > 0$ so, dass für alle $t \in I$ und alle $x_1, x_2 \in U$*

$$\|Y(t, x_1) - Y(t, x_2)\| \leq L \cdot \|x_1 - x_2\| \quad (15.17)$$

Dann existiert für alle $t_0 \in I$, $x_0 \in U$ ein $\delta > 0$ mit $t_0 + \delta \in I$ und so, dass $\forall t_1 \in (t_0, t_0 + \delta]$ eine eindeutige Lösung $x : [t_0, t_1] \rightarrow U$ des AWP

$$\dot{x}(t) = Y(t, x(t)), \quad x(t_0) = x_0 \quad (15.18)$$

existiert.

§ 15. GEWÖHNLICHE DIFFERENTIALGLEICHUNGEN

Beweis. Plan: Wir emulieren den Beweis des Satzes von der Umkehrabbildung 13.1 unter Zuhilfenahme des $\int_a^b$ aus 14.3 als Lösung eines linearen Problems sowie des Banachschen Fixpunktsatzes 4.29 auf einem geeigneten metrischen Raum (von Funktionen $[t_0, t_1] \rightarrow U$).

1. Schritt: (Umschreiben als Fixpunktgleichung) Die Suche nach einer differenzierbaren Funktion $x(\cdot)$ mit

$$\dot{x}(t) = Y(t, x(t)) \quad \text{und} \quad x(t_0) = x_0 \quad (15.19)$$

ist äquivalent zur Suche nach einer stetigen Funktion $x(\cdot)$ mit

$$x(t) = x_0 + \int_{t_0}^t Y(s, x(s)) ds \quad (15.20)$$

(Wegen unseres Hauptsatzes 14.2 ist nämlich $x(\cdot)$ dann schon differenzierbar, wenn sie stetig ist und die Gleichung (15.20) erfüllt.) M.a.W. suchen wir Fixpunkte der Abbildung

$$\begin{aligned} T : \mathcal{C}^0(I, U) &\rightarrow \mathcal{C}^0(I, \mathbb{R}^N) \\ T(x(\cdot))(t) &:= x_0 + \int_{t_0}^t Y(s, x(s)) ds \end{aligned} \quad (15.21)$$

evtl. nach Einschränkung auf geeignete Teilintervalle von I .

2. Schritt: (Vollständige metrische Räume) Für $R > 0$ so, dass $\overline{B_R(x_0)} \subset U$ und $t_1 \in I$, $t_1 > t_0$ betrachten wir

$$X = \{x : [t_0, t_1] \rightarrow \overline{B_R(x_0)} \text{ stetig, } x(t_0) = x_0\} \quad (15.22)$$

mit der von der Supremumsnorm auf $\mathcal{C}^0([t_0, t_1], \mathbb{R}^N)$ induzierten Abstandsfunktion

$$\|x_1(\cdot) - x_2(\cdot)\|_{[t_0, t_1]} := \max\{\|x_1(t) - x_2(t)\|_{\mathbb{R}^N} \mid t \in [t_0, t_1]\} \quad (15.23)$$

(X selbst ist kein Vektorraum!) Gemäss (einer leichten Verallgemeinerung von) 8.5 ist X ein vollständiger metrischer Raum. ($\overline{B_R(x_0)}$ ist als abgeschlossene Teilmenge eines vollständigen normierten Raumes vollständig.) Wir müssen jetzt nur noch unsere beiden Schranken R und δ so wählen, dass falls $t_1 \leq t_0 + \delta$:

(i) $T(X) \subset X$, d.h. T ist eine Selbstabbildung von X .

(ii) T ist kontraktiv, d.h.

$$\|T(x_1(\cdot)) - T(x_2(\cdot))\|_{[t_0, t_1]} \leq c \cdot \|x_1(\cdot) - x_2(\cdot)\|_{[t_0, t_1]} \quad (15.24)$$

für ein $c < 1$. Es stellt sich heraus, dass es genügt, δ anzupassen.

3. Schritt: (Selbstabbildung) Angenommen, $x(\cdot) \in X$, insbesondere $x(t) \in \overline{B_R(x_0)}$ $\forall t \in [t_0, t_1]$. Dann ist

$$\begin{aligned} \|T(x(\cdot))(t) - x_0\| &= \left\| \int_{t_0}^t Y(s, x(s)) ds \right\| \\ &\leq |t - t_0| \cdot \max\{\|Y(s, x(s))\| \mid s \in [t_0, t]\} \end{aligned} \quad (15.25)$$

nach unseren Integralabschätzungen 14.6. Wir fixieren nun $R > 0$ mit $\overline{B_R(x_0)} \subset U$ und $t_* > t_0$ mit $t_* \in I$ und setzen

$$M := \max\{\|Y(t, x)\| \mid t \in [t_0, t_*], x \in \overline{B_R(x_0)}\} < \infty \quad (15.26)$$

denn $[t_0, t_*] \times \overline{B_R(x_0)}$ ist folgenkompakt und stetige Funktionen auf folgenkompakten Mengen sind beschränkt, s. 7.6. Dann ist für $t_1 \leq \max\{t_*, t_0 + \frac{R}{M}\}$ und alle $t \in [t_0, t_1]$ als Folge von (15.25)

$$\|T(x(\cdot))(t) - x_0\| \leq \frac{R}{M} \cdot M = R \quad (15.27)$$

also $T(x) \in X$. (Stetigkeit von $T(x)$ war ja schon klar, ebenso $(t_0) = x_0$.)

4. Schritt (Kontraktivität) Für $x_1(\cdot), x_2(\cdot) \in X$ ist $\forall t \in [t_0, t_1]$

$$\begin{aligned} \|T(x_1(\cdot))(t) - T(x_2(\cdot))(t)\| &= \left\| \int_{t_0}^t (Y(s, x_1(s)) - Y(s, x_2(s))) ds \right\| \\ &\leq |t - t_0| \cdot \max\{\|Y(s, x_1(s)) - Y(s, x_2(s))\| \mid s \in [t_0, t]\} \\ \text{Lipschitz-Bed. (15.17)} &\leq |t - t_0| \cdot L \cdot \max\{\|x_1(s) - x_2(s)\| \mid s \in [t_0, t]\} \\ &\leq \underbrace{|t_1 - t_0| \cdot L}_{=: c} \cdot \|x_1(\cdot) - x_2(\cdot)\|_{[t_0, t_1]} \end{aligned} \quad (15.28)$$

Ist dann t_1 ausserdem noch $< t_0 + \frac{1}{L}$, so folgt $c < 1$, und dann ist T eine Kontraktion.

5. Schritt: (Synthese) Wir sehen: Mit

$$0 < \delta < \max\left\{t_* - t_0, \frac{R}{M}, \frac{1}{L}\right\} \quad (15.29)$$

ist für alle $t_1 \in (t_0, t_0 + \delta]$ T eine Kontraktion. Da die konstante Abbildung $x(t) = x_0$ $\forall t \in [t_0, t_1]$ auf jeden Fall in X liegt, ist X nicht-leer, und wir können mit 4.29 schliessen. $\square$

Beispiel 15.3. Im linearen AWP (15.15) ist $Y(t, x) = Ax$. Ausgehend von der Konstanten $x_0(t) = x_0 \forall t$ erhalten wir durch Iteration von T :

$$\begin{aligned} x_1(t) &= T(x_0(\cdot))(t) = x_0 + \int_0^t Ax_0 ds = x_0 + tAx_0 \\ x_2(t) &= T(x_1(\cdot))(t) = x_0 + \int_0^t A(x_0 + sAx_0) ds = x_0 + tAx_0 + \frac{1}{2}t^2A^2x_0 \\ &\vdots \quad (\text{vollständige Induktion}) \quad \vdots \\ x_n(t) &= \sum_{k=0}^n \frac{t^k}{k!} A^k \end{aligned} \quad (15.30)$$

was für $n \rightarrow \infty$ gegen die Lösung $\exp(tA)x_0$ konvergiert.