

Analysis für Informatik

Vorlesungsskriptum

Institut für Integrierte Schaltungen und Quantum Computing
Johannes Kepler Universität Linz, Österreich

Wintersemester 2025/26

Johannes Kofler



Copyright © 2025/26 Johannes Kofler

Zitieren als:

Johannes Kofler, *Analysis für Informatik Vorlesungsskriptum*, Johannes Kepler Universität Linz, Österreich, 2025/26.

Dieses Skriptum verwendet das Legrand Orange Book L^AT_EX Template (<https://www.latextemplates.com/template/legrand-orange-book>), lizenziert unter CC BY-NC-SA 4.0 (<https://creativecommons.org/licenses/by-nc-sa/4.0/>).

Das Titelbild wurde mit dem Microsoft Designer Image Creator generiert.

Inhaltsverzeichnis

Vorwort	5
1 Grundlagen	7
1.1 Logik	7
1.2 Mengenlehre	10
1.3 Beweisverfahren	12
1.4 Funktionen	13
1.5 Abzählbarkeit	16
1.6 Berechenbarkeit	20
2 Reelle und komplexe Zahlen	25
2.1 Reelle Zahlen	25
2.2 Komplexe Zahlen	29
3 Folgen und Reihen	33
3.1 Folgen	33
3.2 Reihen	42
4 Stetigkeit	53
5 Differentialrechnung	63
5.1 In einer Variablen	63
5.2 In mehreren Variablen	78

6	Integralrechnung	83
6.1	In einer Variablen	83
6.2	In mehreren Variablen	89
6.3	Numerische Integration	93
7	Differentialgleichungen	97
7.1	Gewöhnliche Differentialgleichungen	97
7.2	Partielle Differentialgleichungen	100
7.3	Systeme von Differentialgleichungen	102
7.4	Numerische Verfahren	102
A	Formelsammlung	105

Vorwort

Die Analysis ist jenes Teilgebiet der Mathematik, das sich mit Funktionen und deren Stetigkeit, Differenzierbarkeit und Integrierbarkeit beschäftigt. Die Grundlagen der Analysis wurden durch Gottfried Wilhelm Leibniz und Isaac Newton im 17. Jahrhundert gelegt, die unabhängig voneinander die Infinitesimalrechnung entwickelten (Abbildung 1).

Die Analysis bezog im Laufe ihrer historischen Entwicklung viele ihrer Fragestellungen und Impulse aus den Natur- und Ingenieurwissenschaften und ist dort bis heute von zentraler Bedeutung. Die praktischen Anwendungen sind unzählig und vielfältig und reichen von der Berechnung der Planetenbahnen um die Sonne und der Elektronenorbitale um den Atomkern über elektrische Schaltkreise und Halbleiter bis hin zur mathematischen Beschreibung von computergraphischen Verfahren und künstlichen neuronalen Netzwerken.

Das vorliegende Skriptum entstand weitestgehend im Laufe der ersten Abhaltung der Vorlesung “Analysis für Informatik” im Wintersemester 2024/25. Es werden weiterhin kontinuierlich Verbesserungen eingearbeitet.

Johannes Kofler

Institut für Integrierte Schaltungen und Quantum Computing
Johannes Kepler Universität Linz, Österreich
Wintersemester 2025/26



Abbildung 1: Gottfried Wilhelm Leibniz (1646–1716) und Isaac Newton (1643–1727). Bildquelle: Wikipedia.

1. Grundlagen

Alles ist Zahl.

Pythagoras von Samos (570–510 v.u.Z.)

1.1 Logik

Aussagenlogik

Die Aussagenlogik beschäftigt sich mit Aussagen und deren Verknüpfungen. In der klassischen Aussagenlogik gilt das *Prinzip der Zweiwertigkeit*, und es wird jeder Aussage ein Wahrheitswert der Booleschen Algebra (“wahr” bzw. 1 oder “falsch” bzw. 0) zugeordnet.

Aus atomaren, dh. nicht zusammengesetzten, Aussagen P und Q kann man durch Junktoren zusammengesetzte Aussagen bilden.

- Negation: $\neg P$ (“Nicht P .”)
- Implikation: $P \rightarrow Q$ (“Wenn P , dann Q .”, “ P ist hinreichend für Q .”, “ Q ist notwendig für P .”)
- Bikonditional oder Äquivalenz: $P \leftrightarrow Q$ (“ Q genau dann, wenn P .”, “ P ist notwendig und hinreichend für Q .”)
- Konjunktion: $P \wedge Q$ (“Sowohl P als auch Q .”)
- Disjunktion: $P \vee Q$ (“Entweder P oder Q oder beide.”)

In der klassischen Logik gilt ferner das *Prinzip der Extensionalität*: Aus den Wahrheitswerten von Teilaussagen lässt sich der Wahrheitswert einer zusammengesetzten Aussage eindeutig bestimmen. Beispielsweise ist die Aussage $P \wedge Q \wedge R$ genau dann wahr, wenn alle drei Teilaussagen P , Q , R wahr sind.

Weitere Junktoren lassen sich durch Kombination der obigen Junktoren bilden. Beispielsweise kann das Exklusiv-Oder $\dot{\vee}$ aus Negation, Konjunktion und Disjunktion wie folgt gebaut werden: $P \dot{\vee} Q$ ist äquivalent zu $(\neg P \wedge Q) \vee (P \wedge \neg Q)$.

Innerhalb der klassischen Logik ist folgender Satz herleitbar:¹

¹Der Satz vom ausgeschlossenen Dritten gilt auch in manchen (aber nicht allen) Mehrwertlogiken, in denen es mehr als zwei Wahrheitswerte gibt. Betrachten wir eine Dreiwert-Logik mit den Wahrheitswerten 0

Satz vom ausgeschlossenen Dritten (tertium non datur): Für jede zulässige mathematische Aussage ist entweder die Aussage selbst oder ihr Gegenteil wahr.

Das heißt, für jede Aussage P ist $P \vee \neg P$ eine wahre Aussage. Es gibt neben P und $\neg P$ keine dritte Möglichkeit.

Aus dem Satz vom ausgeschlossenen Dritten folgt mithilfe von Schlussregeln und dem Gesetz der doppelten Negation ($\neg\neg P = P$) der

Satz vom Widerspruch: Für jede beliebige Aussage P gilt $\neg(P \wedge \neg P)$, dh. es ist nicht möglich, dass eine Aussage und ihr Gegenteil gleichzeitig gelten.

Implikationen sind transitiv, dh.:

$$(P \rightarrow Q) \wedge (Q \rightarrow R) \rightarrow (P \rightarrow R) \quad (1.1)$$

Eine *wahre* Implikation $P \rightarrow Q$ bezeichnet man als Folgerung und verwendet die Notation $P \Rightarrow Q$. Eine Kette von Folgerungen

$$A \Rightarrow P_1 \Rightarrow P_2 \Rightarrow \dots P_n \Rightarrow S \quad (1.2)$$

ist ein mathematischer Beweis des Satzes S aus der Annahme A . Eine Äquivalenz “ $\leftrightarrow$ ”, die wahr ist, schreiben wir als “ $\Leftrightarrow$ ”.

Beispiel: Man zeige mittels Wahrheitstabelle (1 = wahr, 0 = falsch) die Äquivalenz $(P \rightarrow Q) \Leftrightarrow (\neg P \vee Q)$. Lösung:

P	Q	$\neg P$	$\neg P \vee Q$	$P \rightarrow Q$
1	1	0	1	1
1	0	0	0	0
0	1	1	1	1
0	0	1	1	1

Dass $0 \rightarrow 0$ und $0 \rightarrow 1$ gilt, bezeichnet man als “*ex falso quodlibet*” (“aus Falschem folgt Beliebiges”).

Prädikatenlogik

Die Prädikatenlogik ist eine Erweiterung der Aussagenlogik, in der zusammengesetzte Aussagen daraufhin untersucht werden können, aus welchen einfacheren Aussagen sie aufgebaut sind. In der Aussage “Sokrates ist ein Mensch” ist “Sokrates” der Eigenname bzw. die Variable und “ist ein Mensch” ist das (einstellige) Prädikat. In der (falschen) Aussage “Sokrates war ein Schüler von Platon” ist “war ein Schüler von” ein zweistelliges Prädikat, weil es eine Relation von zwei Variablen beschreibt.

(falsch), $\frac{1}{2}$ (unbestimmt), und 1 (wahr). Fall (i): Man definiert die Negationen $\neg 0 = 1$, $\neg \frac{1}{2} = \frac{1}{2}$, $\neg 1 = 0$. Für $P = \frac{1}{2}$ ist $P \vee \neg P$ nicht wahr, also ist “tertium non datur” nicht erfüllt. Fall (ii): Man definiert die Negationen $\neg 0 = \frac{1}{2} \vee 1$, $\neg \frac{1}{2} = 0 \vee 1$, $\neg 1 = 0 \vee \frac{1}{2}$. Nun gilt der Satz vom ausgeschlossenen Dritten.

Ein wesentliches Konzept in der Prädikatenlogik sind Quantoren, die angeben, von wie vielen Variablen ein Prädikat erfüllt wird. Die zwei wichtigsten Quantoren sind:

- Der Existenzquantor $\exists$ (“es existiert”): Die Zeichenfolge $\exists n : A(n)$ bedeutet “Es existiert (mindestens ein) n , für das $A(n)$ gilt.”
- Der Allquantor $\forall$ (“für alle”): Die Zeichenfolge $\forall n : A(n)$ bedeutet “Für alle n gilt $A(n)$.”

Beispiel: Die Aussage $\exists n : n + 1 = 2$ ist wahr, weil $n = 1$ die Gleichung erfüllt. Die Aussage $\forall n : n + 1 = 2$ ist falsch, weil die Gleichung nicht für alle n erfüllt ist; $n = 0$ ist ein Gegenbeispiel.

Für Verneinungen von Quantor-Aussagen gelten folgende Äquivalenzen:

$$\neg(\forall n : A(n)) \Leftrightarrow \exists n : \neg A(n), \quad (1.3)$$

$$\neg(\exists n : A(n)) \Leftrightarrow \forall n : \neg A(n). \quad (1.4)$$

In Worten: $A(n)$ gilt nicht für alle n genau dann, wenn es mindestens ein n gibt, sodass $A(n)$ falsch ist. Ferner: Es gibt kein einziges n , sodass $A(n)$ gilt, genau dann, wenn für alle n gilt, dass $A(n)$ falsch ist.

Unentscheidbarkeit

Nicht alle Aussagen sind formal beweis- oder widerlegbar. Die rückbezügliche Aussage “Diese Aussage ist falsch.” ist weder wahr (denn dann wäre sie falsch) noch falsch (denn dann wäre sie wahr). Kurt Gödel bewies 1931 folgendes:

Erster Gödelscher Unvollständigkeitssatz: Jedes hinreichend mächtige, rekursiv aufzählbare formale System ist entweder widersprüchlich oder unvollständig.

In einem formalen System lassen sich mathematische Aussagen aus einer Menge von Schlussregeln aus Axiomen und bereits bewiesenen Aussagen herleiten. Ein formales System ist hinreichend mächtig (oder ausdrucksstark), wenn es in der Lage ist, grundlegende arithmetische Aussagen über die natürlichen Zahlen zu formulieren und zu beweisen, insbesondere Aussagen über Addition und Multiplikation mithilfe von Quantoren. Rekursiv aufzählbar bedeutet, dass alle Aussagen des Systems mithilfe der natürlichen Zahlen aufgezählt (durchnummeriert) werden können.² Widersprüchlich (= inkonsistent) ist ein formales System, wenn es darin Aussagen gibt, von denen sowohl deren Wahrheit als auch deren Falschheit beweisbar ist. Unvollständig ist ein formales System, wenn es darin Aussagen gibt, die weder beweisbar noch widerlegbar sind.³

²Die Arithmetik ist rekursiv aufzählbar, weil sie sich nur mit endlichen Symbolketten über einem abzählbaren Alphabet beschäftigt. Es gibt zwar unendlich viele Aussagen, aber sie können abgezählt (durchnummeriert) werden, selbst wenn diese Aussagen die reellen Zahlen betreffen, die selbst nicht abzählbar sind. Mehr zum Thema Abzählbarkeit dann später.

³Der *zweite Gödelsche Unvollständigkeitssatz* besagt: Jedes hinreichend mächtige und konsistente (= widerspruchsfreie) formale System kann die eigene Konsistenz nicht beweisen.

1.2 Mengenlehre

Es ist eine der grundlegendsten Abstraktionseigenschaften des menschlichen Gehirns, dass wir Objekte nach bestimmten Charakteristika zusammenfassen können. Georg Cantor hat 1895 den mathematischen Begriff der *Menge* eingeführt:

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

In einer Menge kommt jedes Element höchstens einmal vor. Wenn ein Element a in einer Menge M enthalten ist, schreiben wir $a \in M$, andernfalls $a \notin M$. Zwei Mengen A und B sind gleich, dh. $A = B$, wenn sie genau die gleichen Elemente enthalten. Die Reihenfolge spielt keine Rolle. *Beispiel:* Gegeben $A = \{1, 2, 3\}$ und $B = \{3, 2, 1\}$, dann gilt $A = B$. Die leere Menge enthält keine Elemente und wird geschrieben als $\emptyset := \{\}$.

Menge A ist eine Teil- oder Untermenge der Obermenge B , dh. $A \subseteq B$, wenn jedes Element von A auch in B vorkommt. Wenn zusätzlich $A \neq B$ gilt, dann ist A eine echte Teilmenge, und wir schreiben $A \subset B$. *Beispiel:* Gegeben $A = \{2, 3, 5, 7\}$ und $B = \{2, 3, 5, 7, 11, 13, 17\}$, dann gilt $A \subseteq B$ und $A \subset B$.

Mengenoperationen

Mengenoperationen verknüpfen Mengen zu neuen Mengen:

Die Vereinigungsmenge zweier Mengen A und B , geschrieben als $A \cup B$, enthält alle Elemente, die zu A oder B gehören.

Die Schnittmenge (oder der Durchschnitt) $A \cap B$ besteht aus jenen Elementen, die sowohl in A als auch in B sind.

Die Differenzmenge (oder das Komplement) $A \setminus B$ enthält alle Element aus A , die nicht auch in B sind.

Beispiel: Gegeben $A = \{0, 1, 2\}$ und $B = \{1, 2, 3\}$, dann ist $A \cup B = \{0, 1, 2, 3\}$, $A \cap B = \{1, 2\}$, $A \setminus B = \{0\}$.

Falls die Schnittmenge zweier Mengen leer ist, dh. $A \cap B = \emptyset$, dann bezeichnet man diese beiden Mengen als disjunkt.

Potenzmengen

Die Potenzmenge $\mathcal{P}(M)$ einer Menge M , oft auch als 2^M geschrieben, ist die Menge aller Teilmengen von M , dh. $\mathcal{P}(M) := \{U \mid U \subseteq M\}$. Der vertikale Strich wird gelesen als “für die gilt” oder “mit der Eigenschaft”.

Beispiel: Es sei $M = \{a, b\}$. Dann ist $\mathcal{P}(M) = \{\emptyset, \{a\}, \{b\}, \{a, b\}\}$.

Zahlbereiche

Wichtige und häufig verwendete Zahlbereiche sind

- die natürlichen Zahlen $\mathbb{N} := \{1, 2, 3, \dots\}$,
- die natürlichen Zahlen einschließlich des Elements 0: $\mathbb{N}_0 = \mathbb{N} \cup \{0\} = \{0, 1, 2, 3, \dots\}$
- die ganzen Zahlen $\mathbb{Z} := \{\dots, -2, -1, 0, 1, 2, 3, \dots\}$,

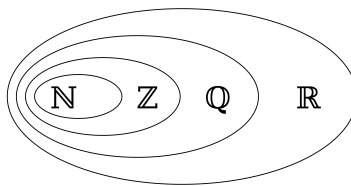


Abbildung 1.1: Mengendiagramm: Die reellen Zahlen $\mathbb{R}$ sind eine echte Obermenge der rationalen Zahlen $\mathbb{Q}$, die wiederum eine echte Obermenge der ganzen Zahlen $\mathbb{Z}$ ist, die wiederum eine echte Obermenge der natürlichen Zahlen $\mathbb{N}$ darstellt. Bildquelle: Wikipedia.

- die rationalen Zahlen $\mathbb{Q} := \{\frac{a}{b} | a \in \mathbb{Z}, b \in \mathbb{N}\}$,
- die reellen Zahlen $\mathbb{R}$, dh. die rationalen und die irrationalen Zahlen,
- die komplexen Zahlen $\mathbb{C} := \{a + ib | a, b \in \mathbb{R}\}$ mit $i := \sqrt{-1}$.

Die Zeichenfolge $:=$ steht für “ist definitionsgemäß gleich”, wobei der Doppelpunkt immer auf der Seite des zu definierenden Symbols steht. Es gilt die Inklusionskette (siehe Abbildung 1.1):

$$\emptyset \subset \mathbb{N} \subset \mathbb{N}_0 \subset \mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R} \subset \mathbb{C}. \quad (1.5)$$

Das Russellsche Paradoxon

Die Cantorsche Mengendefinition schließt nicht aus, dass sich eine Menge selbst enthält. Bertrand Russell hat 1903 gezeigt, dass dies zu Widersprüchen (Antinomien) führt und damit eine tiefe Krise in der Mathematik ausgelöst. Im Wesentlichen lautet Russells Argumentation wie folgt:

Zunächst bezeichnen wir Mengen, die sich nicht selbst enthalten als “normal”, wohingegen wir Mengen, die sich selbst enthalten “nicht normal” nennen. Die Menge $A = \{1, 2, 3\}$ ist normal (da $A \notin A$), wohingegen $B = \{1, 2, 3, B\}$ nicht normal ist (da $B \in B$). Nun definieren wir die Menge M , die alle normalen Mengen enthält und stellen uns die Frage, ob M selbst normal ist. Wäre M normal (dh. $M \notin M$), dann muss M aufgrund seiner Normalität und aufgrund der Definition von M in M liegen (dh. $M \in M$). Also $M \notin M \Rightarrow M \in M$. Wäre M hingegen nicht normal (dh. $M \in M$), dann darf M wegen seiner Nicht-Normalität nicht in M liegen (dh. $M \notin M$). Also $M \in M \Rightarrow M \notin M$. Insgesamt erhalten wir den Widerspruch

$$M \in M \Leftrightarrow M \notin M. \quad (1.6)$$

In Worten: M liegt genau dann in M , wenn M nicht in M liegt.⁴ Diese Antinomie kann nur dadurch verhindert werden, dass man den Mengenbegriff abändert und sich selbst enthaltende Mengen nicht zulässt.

⁴Ein anschauliches Beispiel von Russells Antinomie ist das Barbier-Paradoxon: In einem Dorf gibt es einen Barbier, der all jene, und nur jene, Menschen rasiert, die sich nicht selbst rasieren. Rasiert sich der Barbier selbst oder nicht?

1.3 Beweisverfahren

Ein mathematischer Beweis ist die fehlerfreie Herleitung der Wahrheit oder Falschheit einer Aussage (eines “Satzes”) aus einer Menge von Axiomen (Grundprinzipien) und Prämissen (bereits bewiesene Aussagen). Es gibt mehrere Möglichkeiten, mathematische Beweise durchzuführen. Im folgenden werden drei häufig verwendete Verfahren explizit vorgestellt.

Direkter Beweis

Im direkten Beweis leitet man den zu beweisenden Satz durch logische Schlussfolgerungen aus Prämissen ab. Zur Veranschaulichung betrachten wir das nachstehende einfache Beispiel:

Behauptung: Das Quadrat einer ungeraden natürlichen Zahl ist stets ungerade.
Direkter Beweis: Es sei n eine ungerade natürliche Zahl. Diese kann dargestellt werden als $n = 2k - 1$, wobei k eine beliebige (gerade oder ungerade) natürliche Zahl ist. Quadrieren von n ergibt $n^2 = (2k - 1)^2 = 4k^2 - 4k + 1 = 4k(k - 1) + 1$. Das ist immer eine ungerade Zahl. Was zu beweisen war (*quod erat demonstrandum*).

Das Ende eines Beweises wird oft durch die Buchstaben “q.e.d.” oder ein Quadrat-Symbol “□” dargestellt.

Beweis durch Widerspruch

Ein Widerspruchsbeweis (*reductio ad absurdum*) ist ein indirekter Beweis, bei dem man zeigt, dass ein Widerspruch entsteht, wenn die zu beweisende Aussage falsch wäre. Aus dem Satz vom ausgeschlossenen Dritten folgt dann die Wahrheit der zu beweisenden Aussage. Wir betrachten das folgende Beispiel:⁵

Behauptung: Die Zahl $\sqrt{2}$ ist irrational.
Beweis durch Widerspruch: Nehmen wir an, $\sqrt{2}$ wäre rational, also die Behauptung wäre falsch. Dann kann man $\sqrt{2}$ als Bruch schreiben: $\sqrt{2} = \frac{a}{b}$. Wir dürfen ohne Verlust der Allgemeinheit annehmen, dass a und b teilerfremd sind, dass also der Bruch schon bestmöglich gekürzt wurde. Durch Quadrieren folgt $2 = \frac{a^2}{b^2}$, dh. $a^2 = 2b^2$. Daraus folgt, dass a^2 eine gerade Zahl ist. Da das Quadrat aller ungeraden Zahlen selbst ungerade ist (siehe Beweis oberhalb), muss a gerade sein. Also ist $\frac{a}{2}$ eine natürliche Zahl. Für b^2 können wir schreiben: $b^2 = \frac{a^2}{2} = 2(\frac{a}{2})^2$. Damit ist b^2 gerade, und somit auch b selbst. Da sowohl a als auch b gerade sind, haben sie beide den Teiler 2, was im Widerspruch zur

⁵Der Legende nach löste die Inkommensurabilität von $\sqrt{2}$ eine tiefe Krise der damaligen Mathematik des 5. Jahrhunderts v.u.Z. aus, da die Pythagoräer davon ausgingen, dass nur rationale Zahlen existieren. Hippasos von Metapont, der Entdecker der irrationalen Zahlen, wurde angeblich als Häretiker auf einer Schiffsreise des Pythagoräischen Ordens über Bord geworfen. Wahrscheinlicher ist, dass er aufgrund eines Geheimnisverrats ausgeschlossen wurde und später ertrunken ist. Aber auch diese Legende ist umstritten, da das griechische Wort *árrhētos* für “unsagbar”, das in der Mathematik mit der Bedeutung “irrational” verwendet wurde, auch “geheim” bedeuten konnte.

unproblematischen Annahme der Teilerfremdheit steht. Daher muss die ursprüngliche Annahme, $\sqrt{2}$ wäre rational, falsch sein. Also ist $\sqrt{2}$ irrational. $\square$

Beweis durch vollständige Induktion

Der Beweis einer Aussage durch vollständige Induktion wird oft verwendet um Aussagen der Gestalt “Für jede natürliche Zahl $n \geq n_0$ gilt die Aussage $A(n)$.” zu zeigen. Er basiert auf zwei Schritten. Im *Induktionsanfang* zeigt man zunächst, dass die zu beweisende Aussage für den bestimmten Anfangswert n_0 gilt. Im *Induktionsschritt* zeigt man, dass, wenn die Aussage für n gilt (Induktionsvoraussetzung), sie auch für $n + 1$ wahr ist (Induktionsbehauptung). Wir betrachten wieder ein Beispiel:

Behauptung: Für alle $n \in \mathbb{N}$ gilt $\sum_{k=1}^n (2k - 1) = n^2$.

Beweis durch vollständige Induktion:

- *Induktionsanfang:* Für $n = 1$ ist $\sum_{k=1}^1 (2k - 1) = 2 \cdot 1 - 1 = 1 = 1^2$. Also stimmt die Behauptung für $n = 1$.
- *Induktionsvoraussetzung:* Sei für ein festes, aber beliebiges $n \in \mathbb{N}$ bereits angenommen, dass $\sum_{k=1}^n (2k - 1) = n^2$.
- *Induktionsschritt:* Dann ergibt sich für $n + 1$: $\sum_{k=1}^{n+1} (2k - 1) = \sum_{k=1}^n (2k - 1) + 2(n + 1) - 1 \stackrel{\text{IV}}{=} n^2 + 2n + 1 = (n + 1)^2$, wobei “IV” für Induktionsvoraussetzung steht. Somit ist die Aussage auch für $n + 1$ erfüllt.

Da Induktionsanfang und Induktionsschritt gezeigt sind, gilt die Behauptung für alle $n \in \mathbb{N}$. $\square$

1.4 Funktionen

Das Konzept der Funktion reicht historisch zurück bis in die Antike, wobei die erste explizite Definition im 14. Jahrhundert stattfand. In moderner Ausdrucksweise:

Definition: Eine Funktion (oder Abbildung)

$$f : X \rightarrow Y, x \mapsto y \quad (1.7)$$

ist eine Relation zwischen einer Definitionsmenge X und einer Ziel- oder Bildmenge Y , die jedem Element von X , genannt Funktionsargument oder Urbild x , genau ein Element aus Y , genannt Funktionswert oder Bild $y = f(x)$, zuordnet.

Wir bezeichnen eine Funktion f als

- *surjektiv*, wenn jedes Element y der Bildmenge Y erreicht wird, also mindestens ein Urbild hat: $\forall y \in Y \exists x \in X : f(x) = y$,
- *injektiv*, wenn jedes Bildelement y höchstens einmal erreicht wird, also höchstens ein Urbild hat: $\forall x_1, x_2 \in X : f(x_1) = f(x_2) \Rightarrow x_1 = x_2$,
- *bijektiv*, falls jedes Bildelement y genau ein Urbild hat, falls also f sowohl surjektiv als auch injektiv ist.

Abbildung 1.2 veranschaulicht diese drei Funktionsarten.

Wir betrachten nun den Fall, dass sowohl die Definitions- als auch die Bildmenge

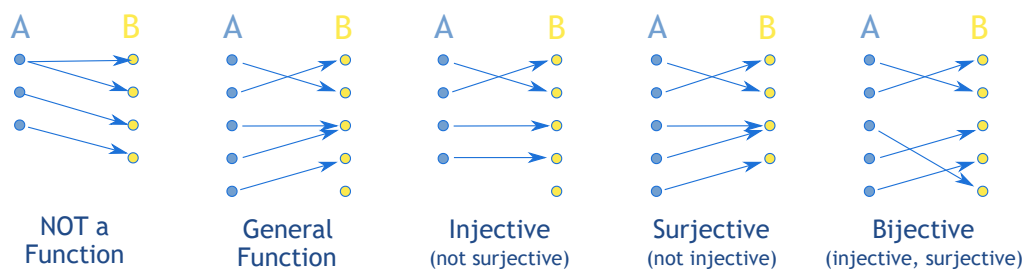


Abbildung 1.2: Von links: (1) Keine Funktion, da ein Element der Definitionsmenge A auf verschiedene Elemente der Bildmenge B verweist. (2) Allgemeine Funktion, weder injektiv noch surjektiv. (3) Injektiv: Jedes Element der Bildmenge wird höchstens einmal erreicht. (4): Surjektiv: Jedes Element der Bildmenge wird mindestens einmal erreicht. (5) Bijektiv: Jedes Element der Bildmenge wird genau einmal erreicht. Bildquelle: *hier*.

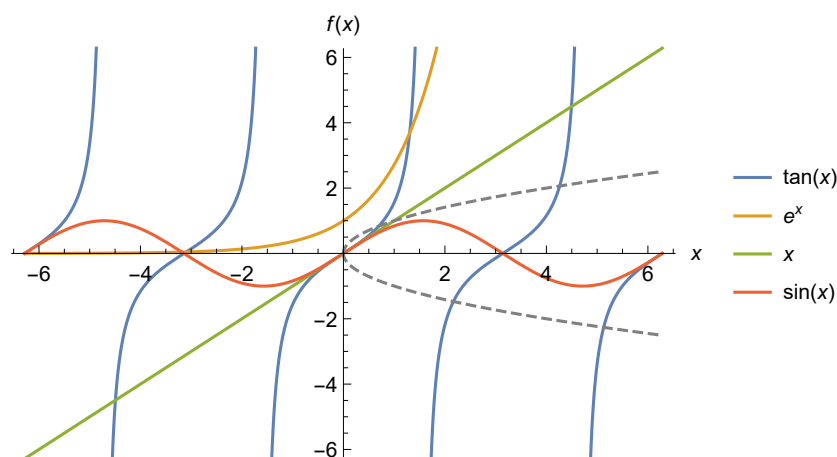


Abbildung 1.3: Gegeben $X = Y = \mathbb{R}$. Die Funktion $f(x) = \tan(x)$ ist surjektiv aber nicht injektiv (blau), $f(x) = e^x$ ist injektiv aber nicht surjektiv (orange), $f(x) = x$ ist bijektiv (grün), und $f(x) = \sin(x)$ ist weder surjektiv noch injektiv (rot). Die "rotierte Parabel" (grau gestrichelt) ist *keine* Funktion, da es x -Werte gibt, denen *mehrere* y -Werte zugeordnet werden.

die reellen Zahlen sind, dh. $X = Y = \mathbb{R}$. Die Funktion $y = \tan(x)$ ist surjektiv (aber nicht injektiv), da alle Bildwerte erreicht werden (und zwar mehrmals). Die Funktion $y = e^x$ ist injektiv (aber nicht surjektiv), da alle Bildwerte höchstens einmal erreicht werden (manche aber gar nicht). Die Funktion $y = x$ ist bijektiv, da jeder Bildwert genau einmal erreicht wird. Die Funktion $f(x) = \sin(x)$ ist weder injektiv noch surjektiv. Die "rotierte Parabel" $y = \pm\sqrt{x}$ ist *keine* Funktion, da jedem (positiven) x -Wert *mehrere* y -Werte zugeordnet werden. Siehe Abbildung 1.3.

Hintereinanderausführung

Zwei Funktionen $f : X \rightarrow Y$ und $g : Y \rightarrow Z$ können hintereinander ausgeführt (verkettet) werden. Wenn zuerst f und dann g ausgeführt wird, ist die resultierende Funktion

$$F := g \circ f : X \rightarrow Z, x \mapsto g(f(x)). \quad (1.8)$$

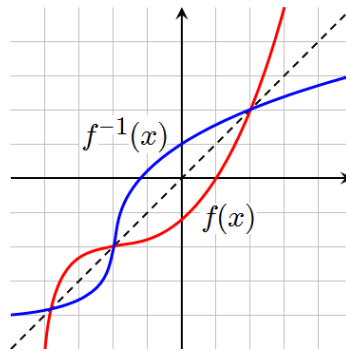


Abbildung 1.4: Graphen der Funktion $f(x)$ (rot) und ihrer Umkehrfunktion $f^{-1}(x)$ (blau). Die gestrichelte Linie ist die Diagonale $y = x$, an der die Funktionen gespiegelt sind. Bildquelle: Wikipedia.

Graphisch sieht das wie folgt aus:

$$\begin{aligned} X &\xrightarrow{f} Y \xrightarrow{g} Z \\ X &\xrightarrow{F=g \circ f} Z \end{aligned} \tag{1.9}$$

Beispiel: Es seien $f(x) = x^2$ und $g(x) = 2x$. Wie lautet die Hintereinanderausführung $F(x) = g(f(x))$?

Lösung: $F(x) = g(f(x)) = g(x^2) = 2x^2$. Man beachte, dass die Reihenfolge wichtig ist: $f(g(x)) = 4x^2$ führt zu einer anderen Funktion, nämlich $G = f \circ g$.

Umkehrfunktion

Für bijektive Funktionen $f : X \rightarrow Y$ kann man eine Umkehrfunktion $g : Y \rightarrow X$ definieren, sodass für jedes x gilt: $g(f(x)) = x$. Anders formuliert ist $g \circ f = \text{id}_X$ die identische Abbildung, die jedem Element x der Menge X sich selbst zuordnet. Die Umkehrfunktion ("Inverse") wird oftmals mit f^{-1} bezeichnet, was in diesem Zusammenhang nicht mit $1/f$ verwechselt werden darf. Die Umkehrfunktion erhält man durch Spiegelung an der Diagonalen $y = x$ (Figur 1.4).

Beispiel: Wie lautet die Umkehrfunktion von $f(x) = 2x + 1$?

Lösung: Aus $f(x) = y = 2x + 1$ folgt $x = (y - 1)/2$. Die Umkehrfunktion von $f(x)$ lautet daher $f^{-1}(x) = (x - 1)/2$.

Probe/Beweis: $f^{-1}(f(x)) = f^{-1}(2x + 1) = (2x + 1 - 1)/2 = x$. $\square$

Wenn f und g bijektive Funktionen sind, dann ist auch $g \circ f$ bijektiv. Die Umkehrfunktion ist $(g \circ f)^{-1} = f^{-1} \circ g^{-1}$.

1.5 Abzählbarkeit

Die Anzahl der Elemente einer Menge M bezeichnet man als *Mächtigkeit* oder *Kardinalität* und schreibt dafür $|M|$. Bei endlichen Mengen genügt es, die Elemente zu zählen. Die Kardinalität bezeichnet man in diesem Fall als “abzählbar endlich”. Beispielsweise hat die leere Menge $\emptyset$ die Mächtigkeit $|\emptyset| = 0$, und die Menge $M = \{2, 4, 6\}$ hat Mächtigkeit $|M| = 3$.

Die natürlichen Zahlen $\mathbb{N}$

Die Mächtigkeit der natürlichen Zahlen $\mathbb{N} := \{1, 2, 3, \dots\}$ ist unendlich. Aber da man die natürlichen Zahlen (mithilfe der natürlichen Zahlen) durchnummerieren, also abzählen kann, nennt man diese Kardinalität “abzählbar unendlich” und verwendet für die Kardinalzahl das Symbol $\aleph_0$:

$$|\mathbb{N}| = \aleph_0. \quad (1.10)$$

$\aleph_0$ die kleinste aller unendlichen Kardinalzahlen, also die kleinste Form der Unendlichkeit.

Wie groß ist die Mächtigkeit der natürlichen Zahlen einschließlich der 0, also von $\mathbb{N}_0 := \{0, 1, 2, 3, \dots\}$? Die naive Antwort wäre, dass $|\mathbb{N}_0| > |\mathbb{N}|$, weil $\mathbb{N}_0$ ja ein Element mehr hat als $\mathbb{N}$. Für endliche Mengen wäre das Argument richtig, aber für unendliche Mengen ist es falsch.

Um dies zu sehen, benötigen wir zunächst eine klare Definition, was es bedeutet, dass zwei Mengen die gleiche Mächtigkeit haben:

Definition: Zwei Mengen A und B sind genau dann *gleichmächtig*, wenn es eine bijektive Abbildung zwischen allen Elementen gibt.

Beispielsweise sind die beiden Mengen $A = \{1, 2, 3\}$ und $B = \{2, 4, 6\}$ gleichmächtig, weil die Funktion $f(a) = 2a$ alle Elemente $a \in A$ auf alle Elemente $b \in B$ abbildet und diese Abbildung bijektiv ist.⁶ Die Äquivalenz

$$\text{Gleichmächtigkeit} \Leftrightarrow \exists \text{ bijektive Abbildung} \quad (1.11)$$

gilt nicht nur für endliche Mengen sondern auch für unendliche.

Nun können wir relativ einfach eine bijektive Abbildung zwischen $\mathbb{N}$ und $\mathbb{N}_0$ finden, nämlich:

$$f: \mathbb{N} \rightarrow \mathbb{N}_0, n \mapsto n - 1. \quad (1.12)$$

Jeder natürlichen Zahl $n \in \mathbb{N}$ wird einfach die nächstkleinere Zahl zugeordnet. So werden alle Elemente aus $\mathbb{N}_0$ genau einmal erreicht. Die beiden Mengen sind also gleichmächtig: $|\mathbb{N}_0| = |\mathbb{N}|$.⁷

⁶Eine weitere Möglichkeit wäre eine bijektive Abbildung in tabellarischer Form, beispielsweise $1 \mapsto 4$, $2 \mapsto 6$, $3 \mapsto 2$.

⁷Auf ähnliche Weise kann man zeigen, dass etwa auch die geraden oder ungeraden Zahlen gleichmächtig zu den natürlichen Zahlen sind.

Die ganzen Zahlen $\mathbb{Z}$

Wie mächtig ist die Menge der ganzen Zahlen $\mathbb{Z} := \{\dots, -2, -1, 0, 1, 2, 3, \dots\}$? Hier wäre die naive Antwort, dass $\mathbb{Z}$ doppelt so unendlich ist wie $\mathbb{N}$. Aber es gibt eine bijektive Abbildung von den natürlichen auf die ganzen Zahlen, $f: \mathbb{N} \rightarrow \mathbb{Z}$, nämlich zB. diese hier:

$n:$	1	2	3	4	5	6	7	8	9	...	n gerade	n ungerade	...
$f(n):$	0	1	-1	2	-2	3	-3	4	-4	...	$\frac{n}{2}$	$-\frac{n-1}{2}$	...

Mithin gilt $|\mathbb{Z}| = |\mathbb{N}|$.

Die rationalen Zahlen $\mathbb{Q}$

Wie verhält es sich mit den rationalen Zahlen $\mathbb{Q} := \{\frac{a}{b} | a \in \mathbb{Z}, b \in \mathbb{N}\}$? Dazu betrachten wir folgende Matrix (Cantors erstes Diagonalargument):

0						
↓						
$+\frac{1}{1}$	→	$+\frac{2}{1}$	$+\frac{3}{1}$	→	$+\frac{4}{1}$	$+\frac{5}{1}$...
	↙		↗	↙		
$-\frac{1}{1}$		$-\frac{2}{1}$	$-\frac{3}{1}$		$-\frac{4}{1}$	$-\frac{5}{1}$...
↓	↗	↙	↗	↙	↗	↙
$+\frac{1}{2}$		$+\frac{2}{2}$	$+\frac{3}{2}$		$+\frac{4}{2}$	$+\frac{5}{2}$...
	↙		↗			
$-\frac{1}{2}$		$-\frac{2}{2}$	$-\frac{3}{2}$		$-\frac{4}{2}$	$-\frac{5}{2}$...
↓	↗	↙	↗	↙	↗	↙
$+\frac{1}{3}$		$+\frac{2}{3}$	$+\frac{3}{3}$		$+\frac{4}{3}$	$+\frac{5}{3}$...
⋮		⋮	⋮		⋮	⋮

Die erste Zeile enthält nur die 0. Die nächste Zeile listet alle positiven Brüche mit Nenner 1 auf. Die nächste alle negativen Brüche mit Nenner 1. Die nächste alle positive Brüche mit Nenner 2. Und so weiter. Die Matrix enthält daher alle rationalen Zahlen. Durch “schlangenförmiges diagonales Abwandern” (rote Pfeile) wird vermieden, in einer Richtung ins Unendliche gehen zu müssen, von wo aus es keinen Weg zurück gäbe. Somit kann man die rationalen Zahlen abzählen, dh. mit den natürlichen Zahlen durchnummerieren:

$$f: \mathbb{N} \rightarrow \mathbb{Q}, 1 \mapsto 0, 2 \mapsto +\frac{1}{1}, 3 \mapsto +\frac{2}{1}, 4 \mapsto -\frac{1}{1}, 5 \mapsto +\frac{1}{2}, 6 \mapsto -\frac{2}{1}, \dots \quad (1.13)$$

Sich wiederholende rationale Zahlen (in grau geschrieben) – die erste ist $+\frac{2}{2}$, da $+\frac{1}{1}$ schon gezählt wurde – werden einfach ausgelassen. Die Menge der rationalen Zahlen ist also abzählbar unendlich. Wir haben mithin

$$|\mathbb{Q}| = |\mathbb{N}|. \quad (1.14)$$

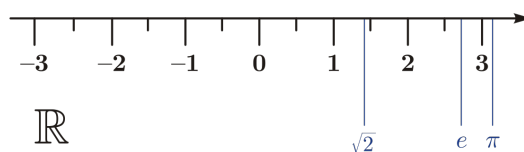


Abbildung 1.5: Die Zahlengerade ist eine Veranschaulichung der reellen Zahlen $\mathbb{R}$. Bildquelle: Wikipedia.

Die reellen Zahlen $\mathbb{R}$

Die reellen Zahlen $\mathbb{R}$ beinhalten auch die irrationalen Zahlen wie $\sqrt{2}$, e und π (Abbildung 1.5). Die rationalen Zahlen liegen *dicht* in den reellen. Das bedeutet, dass für jede reelle Zahl $r \in \mathbb{R}$ und vorgegebene Fehlerschranke $\varepsilon > 0$ eine rationale Zahl $q \in \mathbb{Q}$ gefunden werden kann, sodass $|r - q| < \varepsilon$. Man kann jede reelle Zahl also immer beliebig gut durch eine rationale Zahl approximieren.

Wie sieht es nun mit der Mächtigkeit der reellen Zahlen aus? Die Antwort auf diese Frage gibt der

Satz von Cantor: Die Menge der reellen Zahlen $\mathbb{R}$ ist überabzählbar unendlich.

Beweis (Cantors zweites Diagonalargument): Bewiesen wird im folgenden die noch stärkere Aussage, dass sogar das Intervall I der reellen Zahlen zwischen 0 und 1, $I = [0, 1[= \{x \in \mathbb{R} | 0 \leq x < 1\} \subset \mathbb{R}$ überabzählbar ist.^a Durchgeführt wird ein Beweis durch Widerspruch. Wir nehmen daher an, dass I abzählbar ist. Dann kann man *alle* reellen Zahlen $r \in I$ mit den natürlichen Zahlen $n \in \mathbb{N}$ durchnummerieren, dh. es existiert $f: \mathbb{N} \rightarrow \mathbb{R}$, $n \mapsto r_n$, und in Dezimalschreibweise auflisten:^b

$$\begin{aligned} r_1 &= 0.\underline{r_1^{(1)}} \underline{r_1^{(2)}} r_1^{(3)} r_1^{(4)} r_1^{(5)} \dots \\ r_2 &= 0.\underline{r_2^{(1)}} \underline{r_2^{(2)}} \underline{r_2^{(3)}} r_2^{(4)} r_2^{(5)} \dots \\ r_3 &= 0.\underline{r_3^{(1)}} r_3^{(2)} \underline{r_3^{(3)}} \underline{r_3^{(4)}} r_3^{(5)} \dots \\ &\dots \end{aligned} \tag{1.15}$$

Hier ist $r_n^{(j)} \in \{0, 1, 2, \dots, 9\}$ die j -te Dezimalstelle der n -ten Zahl r_n . Die Diagonalstellen $r_n^{(n)}$ sind unterstrichen. Nun definieren wir eine Zahl

$$s = 0.s^{(1)} s^{(2)} s^{(3)} s^{(4)} s^{(5)} \dots, \tag{1.16}$$

wobei $s^{(n)} := 1$ falls $r_n^{(n)} = 0$ und $s^{(n)} := 0$ falls $r_n^{(n)} \neq 0$. Die Zahl s weicht also von r_1 auf der ersten Dezimalstelle ab, da $s^{(1)} \neq r_1^{(1)}$. Von r_2 weicht s auf der zweiten Dezimalstelle ab, da $s^{(2)} \neq r_2^{(2)}$. Von r_3 auf der dritten, usw. Das bedeutet, dass sich s *nicht* in der Liste der Zahlen $r_1, r_2, r_3, \dots$ befindet. Dies wiederum steht im Widerspruch zu unserer Annahme, dass wir alle reellen Zahlen zwischen 0 und 1 aufgelistet haben. Dieser

Widerspruch lässt sich nur auflösen, indem wir folgern, dass das Abzählverfahren $n \mapsto r_n$ nicht existiert. Die reellen Zahlen sind also überabzählbar unendlich. $\square$

^aEine offene (geschlossene) eckige Klammer bedeutet, dass das Randelement (nicht) Teil der Menge ist.

^bWir verwenden den im Englischen und in der Computerwissenschaft üblichen Dezimalpunkt, nicht das deutsche Dezimalkomma.

Die Mächtigkeit der reellen Zahlen zwischen 0 und 1 ist gleich der Mächtigkeit aller reellen Zahlen. Eine Bijektion, die das zeigt, lautet $f: I \rightarrow \mathbb{R}, x \mapsto \tan[\pi(x + \frac{1}{2})]$.

Die Kardinalzahl der reellen Zahlen wird üblicherweise mit $\mathfrak{c}$ (für das lateinische Wort *continuum*) bezeichnet. Wir haben mithin:

$$|\mathbb{N}| = |\mathbb{Z}| = |\mathbb{Q}| = \aleph_0 < \mathfrak{c} = |\mathbb{R}|. \quad (1.17)$$

Die *Kontinuumshypothese* – aufgestellt von Cantor im Jahr 1878 – besagt, dass es zwischen $\aleph_0$ und $\mathfrak{c}$, also zwischen der Unendlichkeit der natürlichen Zahlen und jener der reellen Zahlen, keine anderen Unendlichkeiten gibt. (Anders formuliert besagt die Kontinuumshypothese, dass $\aleph_1$, definiert als die nächst kleinste Unendlichkeit nach $\aleph_0$, gleich $\mathfrak{c}$ ist.) Gödel bewies 1938, dass sich die Kontinuumshypothese aus der Standard-Mengenlehre⁸ nicht widerlegen lässt. Paul Cohen zeigte 1963 wiederum, dass sich die Kontinuumshypothese aus der Standard-Mengenlehre nicht beweisen lässt. Die Kontinuumshypothese kann also weder bewiesen noch widerlegt werden, ist folglich unentscheidbar. Sie ist damit ein Beispiel für Gödels (ersten) Unvollständigkeitssatz.

Die komplexen Zahlen $\mathbb{C}$ und Vektorräume $\mathbb{R}^N$

Komplexe Zahlen aus $\mathbb{C} := \{a + ib | a, b \in \mathbb{R}\}$ mit $i = \sqrt{-1}$ sind Paare (a, b) von zwei reellen Zahlen und damit gleichmächtig zur euklidischen Ebene $\mathbb{R}^2 = \mathbb{R} \times \mathbb{R}$, in der Vektoren mit zwei Elementen leben. Das Symbol $\times$ bezeichnet das kartesische Produkt zweier Mengen, das die Menge aller geordneten Paare bildet, wobei die erste Komponente ein Element der ersten Menge und die zweite Komponente ein Element der zweiten Menge ist. Vektorräume kann man in beliebig hoher Dimension $N \in \mathbb{N}$ konstruieren mit Elementen $(x_1, x_2, \dots, x_N) \in \mathbb{R}^N = \mathbb{R} \times \mathbb{R} \times \dots \times \mathbb{R}$. In den Natur- und Ingenieurwissenschaften verwendet man üblicherweise Spaltenvektoren $\vec{x} \equiv \mathbf{x} = (x_1, x_2, \dots, x_N)^T$, wobei das T (für “transponiert”) aus einem Zeilenvektor einen Spaltenvektor macht und umgekehrt.

Alle $\mathbb{R}^N$ sind gleichmächtig zu den reellen Zahlen.⁹ Dies führt zu folgender Zusammenfassung:

$$|\mathbb{N}| = |\mathbb{Z}| = |\mathbb{Q}| = \aleph_0 < \mathfrak{c} = |\mathbb{R}| = |\mathbb{C}| = |\mathbb{R}^N|. \quad (1.18)$$

⁸Gemeint sind hier die sogenannten Zermelo-Fraenkel-Axiome inklusive Auswahlaxiom. Eine genauere Diskussion würde hier zu weit führen.

⁹Beweisidee: Zunächst betrachten wir folgende bijektive Abbildung von $[0, 1]^2$ nach $[0, 1]$: Die Dezimalstellen werden verschachtelt, sodass $(r; s) = (0.r^{(1)}r^{(2)}r^{(3)} \dots, 0.s^{(1)}s^{(2)}s^{(3)} \dots)$ abgebildet wird auf $0.r^{(1)}s^{(1)}r^{(2)}s^{(2)}r^{(3)}s^{(3)} \dots$. Dieses Verfahren kann man dann erweitern auf zunächst alle reellen Zahlen, nicht nur jene zwischen 0 und 1. Und dann von 2 auf N Dimensionen.

Potenzmengen

Die Potenzmenge $\mathcal{P}(M)$ einer Menge M hat stets eine größere Kardinalität als M . (Für endliche Mengen ist dies sofort klar. Für unendliche Mengen folgt die Aussage über eine Verallgemeinerung des zweiten Cantorschen Diagonalarguments.) Es gibt also eine unendliche Leiter von Unendlichkeiten:

$$\aleph_0 < \mathfrak{c} = 2^{\aleph_0} < 2^{\mathfrak{c}} < 2^{2^{\mathfrak{c}}} < \dots \quad (1.19)$$

1.6 Berechenbarkeit

Definition: Eine reelle Zahl

$$r = r^{(0)}.r^{(1)}r^{(2)}r^{(3)}r^{(4)}r^{(5)}\dots \text{ mit } r^{(0)} \in \mathbb{Z}, r^{(j)} \in \{0, 1, 2, \dots, 9\}, j \in \mathbb{N} \quad (1.20)$$

heißt *berechenbar*, falls sie durch eine berechenbare Funktion $f: \mathbb{N} \rightarrow \mathbb{Q}$, $n \mapsto f(n)$ mit beliebiger Genauigkeit approximiert werden kann. Konkret, wenn für jedes n der Abstand der reellen Zahl r zur rationalen Näherung $f(n)$ kleiner ist als $\frac{1}{n}$:

$$|r - f(n)| < \frac{1}{n}. \quad (1.21)$$

Die 1936 von Alan Turing eingeführte *Turing-Maschine* ist ein mathematisch abstraktes Modell eines Computers. Alle echten Computerarchitekturen und Computerprogramme (C, Java, Python, Mathematica, etc.) sind spezielle Realisierungen. Die – unbewiesene aber weithin anerkannte – Church-Turing-These besagt, dass alles, was prinzipiell berechenbar ist, von einer Turing-Maschine berechnet werden kann – für jede berechenbare Funktion $f(n)$ gibt es also eine Turing-Maschine M mit Input n , die diese Funktion berechnet.

Alle algebraischen Zahlen¹⁰ (zB. $\sqrt{2}$) sind berechenbar, und auch viele transzendente Zahlen¹¹ (zB. die Eulersche Zahl e und die Kreiszahl π):

$$\sqrt{2} = \sum_{k=0}^{\infty} \frac{(2k+1)!}{2^{3k+1}(k!)^2}, \quad (1.22)$$

$$e = \frac{1}{0!} + \frac{1}{1!} + \frac{1}{2!} + \frac{1}{3!} + \dots = \sum_{k=0}^{\infty} \frac{1}{k!}, \quad (1.23)$$

$$\pi = 4 \left(1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} + \dots \right) = 4 \sum_{k=0}^{\infty} \frac{(-1)^k}{2k+1}. \quad (1.24)$$

Da ein Computer nicht unendlich viele Terme addieren kann, müssen wir “ausführbare” Funktionen definieren, indem nur bis zum Index n summiert wird. Die Approximation für π lautet dann $f(n) = 4 \sum_{k=0}^n \frac{(-1)^k}{2k+1}$.

Bezeichnen wir die Menge aller berechenbaren reellen Zahlen mit $\mathbb{B}$ und die Menge

¹⁰Algebraische Zahlen sind die Nullstellen eines Polynoms mit rationalen Koeffizienten. Beispielsweise ist eine Nullstelle des Polynoms $x^2 - 2$ die algebraische Zahl $\sqrt{2}$.

¹¹Transzendente Zahlen sind Zahlen, die nicht algebraisch sind.

aller Computerprogramme bzw. Turing-Maschinen mit $\mathbb{M}$. Für jede berechenbare Zahl $b \in \mathbb{B}$ existiert (mindestens) eine Turing-Maschine $M \in \mathbb{M}$, die diese Zahl berechnet. Also gilt jedenfalls $|\mathbb{B}| \leq |\mathbb{M}|$.

Satz: Die Menge der berechenbaren Zahlen $\mathbb{B}$ und die Menge der Turing-Maschinen $\mathbb{M}$ sind beide abzählbar unendlich.

Beweis: Jede Turing-Maschine M ist durch ihren Quelltext repräsentiert. Dieser Quelltext kann unabhängig vom gewählten Alphabet eineindeutig – dh. in jede Richtung eindeutig bzw. bijektiv – in Binärform konvertiert werden. Und dieser Binärstring entspricht wiederum eineindeutig einer natürlichen Zahl $m \in \mathbb{N}$. Andererseits kann man *jede* natürliche Zahl m als Computerprogramm-Quelltext auffassen.^a Alle prinzipiell möglichen Computerprogramme können also durchnummeriert werden, mithin gibt es abzählbar unendlich viele und es gilt $|\mathbb{M}| = \aleph_0$. Manche Programme, die voneinander verschieden sind, berechnen trotzdem die gleiche Zahl.^b Daher gibt es *höchstens* abzählbar unendlich viele berechenbare Zahlen, also entweder abzählbar endlich viele oder abzählbar unendlich viele. Es müssen aber in der Tat abzählbar unendlich viele sein, da zB. alle rationalen Zahlen berechenbar sind. Daher gilt:

$$|\mathbb{B}| = |\mathbb{M}| = \aleph_0. \quad (1.25)$$

Die Mengen der berechenbaren Zahlen $\mathbb{B}$ und Turing-Maschinen $\mathbb{M}$ sind also beide abzählbar unendlich. $\square$

^aDie meisten Programme werden “unsinnig” sein und sofort mit Fehler terminieren, da ja nur ganz bestimmte Syntax erlaubt ist. Aber das soll uns hier nicht weiter stören.

^bEs gibt beispielsweise viele verschiedene Formeln und daher Programme zur Berechnung der Zahl π .

Da die reellen Zahlen überabzählbar mächtig sind, sind die allermeisten reellen Zahlen folglich nicht berechenbar. Warum steht dies nicht im Widerspruch dazu, dass die rationalen Zahlen (die ja alle berechenbar sind) die reellen Zahlen dicht approximieren? Der Grund liegt darin, dass – nur weil eine gesuchte reelle Zahl beliebig dicht an einer rationalen Zahl liegt – es nicht notwendigerweise ein Programm gibt, das die gesuchte reelle Zahl berechnet. Wir werden gleich ein berühmtes Beispiel einer nicht berechenbaren Zahl kennenlernen.

Das Halteproblem

Das Halteproblem ist eine fundamentale Frage der theoretischen Informatik und wurde von Turing im Jahr 1937 formuliert. Wir skizzieren das Argument im folgenden nur grob; mehr Details werden im Kurs “Berechenbarkeit und Komplexität” behandelt.

Eine Turing-Maschine M nimmt einen Input-String n (den wir uns als natürliche Zahl denken dürfen, dh. $n \in \mathbb{N}$) und beginnt mit ihrer Arbeit. Dann gibt es zwei Möglichkeiten:

1. M hält irgendwann mit Output $M(n)$.¹²
2. M hält nicht und läuft für immer weiter.

Ob eine Turing-Maschine hält oder nicht, ist manchmal sehr einfach, aber oft auch sehr

¹²Der Output $M(n)$ kann auch eine Fehlernachricht wegen schlechter Programm-Syntax sein.

schwer festzustellen. Wir betrachten ein paar

Beispiele:

- program $M_i(n)$

$n := n+1$

return n

Dieses Programm hält für jeden Input n und gibt die nachfolgende Zahl $n+1$ aus.

- program $M_{ii}(n)$

while true

$n := n+1$

return n

Dieses Programm wird niemals anhalten und ewig weiterzählen.

- program $M_{iii}(n)$

if $n < 10$

loop forever

else

return $1/n$

Dieses Programm läuft für alle Inputs $n < 10$ für immer. Für Inputs $n \geq 10$ hält es und gibt die rationale Zahl $1/n$ aus.

- program $M_{iv}(n)$

$k := 4$

while true

if $k = \text{sum of two prime numbers}$

$k := k+2$

else

return k

Dieses Programm ignoriert den Input n . Es terminiert nur, wenn eine gerade Zahl größer als 2 gefunden wird, die sich *nicht* als Summe zweier Primzahlen schreiben lässt. Hält M_{iv} jemals oder läuft es für immer weiter?

Definition: Das *Halteproblem* ist ein Entscheidungsproblem, nämlich die Frage, ob es für *jede* Turing-Maschine M und *jeden* Input n möglich ist festzustellen, ob $M(n)$ hält oder nicht. Falls ein Algorithmus H existiert, der das für *alle* $M(n)$ leistet, ist das Entscheidungsproblem entscheidbar (bzw. gelöst). Wenn so ein Algorithmus nicht existiert, ist das Halteproblem unentscheidbar.

Satz: Das Halteproblem ist unentscheidbar.

Beweis: Eine *universelle Turing-Maschine* U kann als Input jede andere Turing-Maschine M und deren Input $n \in \mathbb{N}$ akzeptieren. Insbesondere kann U so programmiert sein, dass sie M mit Input n ausführt ("simuliert") und den Wert $M(n)$ als Output ausgibt, oder eine bestimmte Funktion von $M(n)$.

Nehmen wir an, es gäbe eine universelle Turing-Maschine H , die für eine ihr überge-

eine Turing-Maschine M und Input n folgendes ausgibt:

$$H(M, n) = \begin{cases} 1 & \text{falls } M(n) \text{ h\"alt,} \\ 0 & \text{falls } M(n) \text{ nicht h\"alt.} \end{cases} \quad (1.26)$$

Solch eine Maschine H l\"ost das Halteproblem, kann also stets entscheiden, ob eine Maschine M mit Input n h\"alt oder nicht. Die Annahme der Existenz von H hat enorme konzeptionelle und praktische Bedeutung und w\"urde eine Unzahl an offenen Problemen der Mathematik l\"osen.^a

Der zentrale Schritt in Turings Argument ist nun, dass wir eine Maschine D konstruieren k\"onnen, die (i) H als Unterprogramm enth\"alt, (ii) sich selbst an H \"ubergibt und (iii) das Gegenteil von dem macht, was H \"uber D sagt:

$$D(n) = \begin{cases} \text{h\"alt} & \text{falls } H(D, n) = 0, \text{ also falls } D(n) \text{ nicht h\"alt,} \\ \text{h\"alt nicht} & \text{falls } H(D, n) = 1, \text{ also falls } D(n) \text{ h\"alt.} \end{cases} \quad (1.27)$$

$D(n)$ h\"alt, wenn $D(n)$ nicht h\"alt, und $D(n)$ h\"alt nicht, wenn $D(n)$ h\"alt:

$$D(n) \text{ h\"alt} \Leftrightarrow D(n) \text{ h\"alt nicht.} \quad (1.28)$$

Damit landen wir bei einer Antinomie, ganz so wie bei Russells Menge aller normalen Mengen und G\"odels r\"uckbezuglicher Aussage "Diese Aussage ist falsch". Der Widerspruch (1.28) kann nur dadurch aufgel\"ost werden, indem die zun\"achst angenommene Existenz von H zur\"uckgewiesen wird. Das Halteproblem ist mithin unentscheidbar. $\square$

^aDie Goldbachsche Vermutung ist eine ber\"uhmte unbewiesene Aussage der Zahlentheorie aus dem Jahr 1742, die besagt, dass sich jede gerade Zahl, die gr\"o\"oer als 2 ist, als Summe zweier Primzahlen schreiben l\"asst. Mithilfe von H entscheiden zu k\"onnen, ob die weiter oben definierte Maschine M_{iv} h\"alt (bzw. nicht h\"alt), w\"urde die Goldbachsche Vermutung widerlegen (bzw. beweisen).

Wie weiter oben bereits diskutiert, kann man jeder nat\"urlichen Zahl $m \in \mathbb{N}$ eineindeutig eine Turing-Maschine M_m zuordnen. Nun k\"onnen wir wie in Cantors erstem Diagonalargument alle Turing-Maschinen M_m mit allen m\"oglichen Inputs $n \in \mathbb{N}$ wie folgt anordnen:

$$\begin{array}{cccccc} n = 1 : & n = 2 : & n = 3 : & n = 4 : & n = 5 : & \dots \\ m = 1 : & M_1(1) & M_1(2) & M_1(3) & M_1(4) & M_1(5) & \dots \\ m = 2 : & M_2(1) & M_2(2) & M_2(3) & M_2(4) & M_2(5) & \dots \\ m = 3 : & M_3(1) & M_3(2) & M_3(3) & M_3(4) & M_3(5) & \dots \\ \dots & & & & & & \end{array} \quad (1.29)$$

Durch diagonales Abz\"ahlen kann man jeder nat\"urlichen Zahl $k \in \mathbb{N}$ eineindeutig eine Turingmaschine $M_m(n)$ zuordnen: $1 \mapsto M_1(1)$, $2 \mapsto M_1(2)$, $3 \mapsto M_2(1)$, $4 \mapsto M_3(1)$, $5 \mapsto M_2(2)$, usw. Der einfachen Notation halber schreiben wir $T_k := M_m(n)$, wobei jedem Paar (m, n) genau ein k entspricht und umgekehrt.¹³

¹³Man kann sich T_k als eine Turing-Maschine vorstellen, deren Input hartkodiert ist, also in k enthalten ist.

Nun definieren wir die folgende reelle Zahl (in Binärdarstellung):

$$h = 0.h^{(1)}h^{(2)}h^{(3)}h^{(4)}h^{(5)}\dots \quad \text{mit } h^{(k)} = \begin{cases} 1 & \text{falls } T_k \text{ hält,} \\ 0 & \text{falls } T_k \text{ nicht hält.} \end{cases} \quad (1.30)$$

Die k -te Dezimalstelle von h ist 1 (bzw. 0) falls T_k hält (bzw. nicht hält). Wäre h berechenbar, dann wäre das Halteproblem entscheidbar, zumal man für jede Turing-Maschine M_m mit jedem Input n voraussagen könnte, ob sie hält oder nicht. Das Halteproblem ist aber unentscheidbar. Also ist h nicht berechenbar, dh. $h \in \mathbb{R} \setminus \mathbb{B}$.¹⁴

¹⁴Ähnlich wie h definiert ist die Chaitinsche Konstante (oder Haltewahrscheinlichkeit) Ω , die die Wahrscheinlichkeit angibt, dass eine zufällig gewählte Turing-Maschine hält. Auch Ω ist unberechenbar.

2. Reelle und komplexe Zahlen

2.1 Reelle Zahlen

Die Konstruktion der reellen Zahlen als Erweiterung des Zahlenbereichs der rationalen Zahlen war ein wichtiger Schritt der Mathematik des 19. Jahrhunderts, der die Analysis auf ein solides mathematisches Fundament stellte.

Die reellen Zahlen können aber nicht nur über die Mengenlehre als Erweiterung der rationalen Zahlen definiert werden, sondern direkt durch Axiome. Man benötigt dazu drei Gruppen von Axiomen: (i) die Körperaxiome, (ii) die Axiome der Ordnungsstruktur sowie (iii) ein Axiom, das die Vollständigkeit garantiert.

Körperaxiome

Die reellen Zahlen $\mathbb{R}$ bilden einen *Körper*, da die zwei Verknüpfungen “+” (Addition) und “·” (Multiplikation), die zwei reellen Zahlen a und b ihre Summe $a + b$ bzw. ihr Produkt $a \cdot b = ab$ zuordnen, die nachstehenden Körperaxiome erfüllen:

- Kommutativgesetze: $a + b = b + a$ und $ab = ba$
- Assoziativgesetze: $a + (b + c) = (a + b) + c$ und $a(bc) = (ab)c$
- Distributivgesetz: $a(b + c) = ab + ac$
- Existenz neutraler Elemente: $a + 0 = a$ und $a \cdot 1 = a$
- Existenz inverser Elemente: $a + (-a) = 0$ und (für $a \neq 0$): $a \cdot a^{-1} = 1$

Ordnungsaxiome

Ferner gibt es auf $\mathbb{R}$ eine *Ordnung* “ $\leq$ ” mit den Ordnungsaxiomen:

- Reflexivität: $a \leq a$
- Transitivität: $a \leq b \wedge b \leq c \Rightarrow a \leq c$
- Antisymmetrie (Identität): $a \leq b \wedge b \leq a \Rightarrow a = b$
- Totalität: $a \leq b \vee b \leq a$
- Verträglichkeit mit Addition: $a \leq b \Rightarrow a + c \leq b + c$
- Verträglichkeit mit Multiplikation: $a \leq b \wedge c \geq 0 \Rightarrow ac \leq bc$

Ordnungsvollständigkeit

Die reellen Zahlen sind *ordnungsvollständig* mit dem Schnittaxiom (Axiom der Ordnungsvollständigkeit):

- Voraussetzungen/Definitionen: A und B seien nichtleere Teilmengen von $\mathbb{R}$ und es gilt $A \cup B = \mathbb{R}$. Für alle $a \in A$ und $b \in B$ gelte $a < b$. Eine Zahl t heißt Trennungszahl des *Dedekindschen Schnittes* $A|B$, wenn für alle a und b gilt: $a \leq t \leq b$.
- Schnittaxiom: Jeder Dedekindsche Schnitt besitzt genau eine Trennungszahl.

Intervalle

Gegeben zwei reelle Zahlen a und $b > a$, definieren wir die Intervalle (abgeschlossen, linksoffen, rechtsoffen, offen):

$$[a, b] = \{x \in \mathbb{R} | a \leq x \leq b\}, \quad (2.1)$$

$$]a, b] = \{x \in \mathbb{R} | a < x \leq b\}, \quad (2.2)$$

$$[a, b[= \{x \in \mathbb{R} | a \leq x < b\}, \quad (2.3)$$

$$]a, b[= \{x \in \mathbb{R} | a < x < b\}. \quad (2.4)$$

Maximum und Minimum

Für zwei reelle Zahlen a und b , definieren wir das Maximum bzw. Minimum:

$$\max(a, b) := \begin{cases} a & \text{falls } a \geq b, \\ b & \text{sonst.} \end{cases} \quad (2.5)$$

$$\min(a, b) := \begin{cases} a & \text{falls } a \leq b, \\ b & \text{sonst.} \end{cases} \quad (2.6)$$

Für mehrere reelle Zahlen definieren wir:

$$\max(a_1, a_2, \dots, a_n) := \max(a_1(\max(a_2, \max(\dots, \max(a_{n-1}, a_n) \dots))), \quad (2.7)$$

$$\min(a_1, a_2, \dots, a_n) := \min(a_1(\min(a_2, \min(\dots, \min(a_{n-1}, a_n) \dots))). \quad (2.8)$$

Supremum und Infimum

Eine Menge $A \subseteq \mathbb{R}$ hat eine obere Schranke $b \in \mathbb{R}$ falls gilt

$$\exists b \in \mathbb{R} \forall a \in A : a \leq b \quad (2.9)$$

Ein Element b heißt Supremum von A , $\sup A$, wenn es die kleinste obere Schranke von A ist. Ganz analog dazu ist das Infimum von A , $\inf A$, die größte untere Schranke von A .

Beispiel: Wir betrachten das beidseitig geschlossene Intervall $A = [0, 1]$, in dem die Zahlen 0 und 1 enthalten sind. Es gilt $\sup A = 1$ und $\inf A = 0$. Hier sind auch das größte Element der Menge A , das Maximum $\max A$, und das kleinste Element, das Minimum $\min A$ definiert: $\max A = 1$ und $\min A = 0$. Ganz allgemein gilt bei geschlossenen

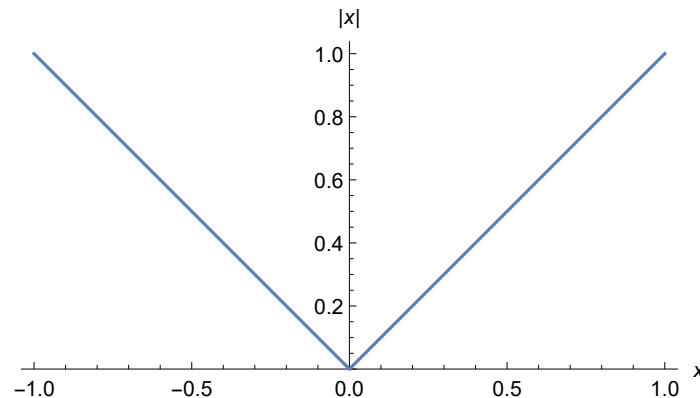


Abbildung 2.1: Die Betragsfunktion $f(x) = |x|$ ordnet jeder Zahl x ihren Abstand zu 0 zu, den sogenannten Absolutbetrag oder Absolutwert.

Intervallen oder Mengen mit endlich vielen Elementen: $\sup A = \max A$ und $\inf A = \min A$.

Beispiel: Wir betrachten das beidseitig offene Intervall $A =]0, 1[$, in dem die Zahlen 0 und 1 nicht enthalten sind. Es gilt $\sup A = 1$ und $\inf A = 0$. Supremum und Infimum sind nicht in A enthalten. Die Menge A hat weder ein größtes Element noch ein kleinstes Element: $\max A$ und $\min A$ sind nicht definiert.

Für nach oben bzw. unten unbeschränkte Mengen A (wie z.B. für $A = \mathbb{R}$) setzen wir $\sup A = \infty$ bzw. $\inf A = -\infty$. Das Unendlichkeitssymbol ∞ steht dabei für *keine* natürliche oder reelle Zahl. Maximum bzw. Minimum sind für unbeschränkte Mengen nicht definiert.

Beispiel: Die positiven reellen Zahlen inklusive der Zahl 0, $\mathbb{R}_0^+ = \{x \in \mathbb{R} \mid x \geq 0\}$, sind ein nach links abgeschlossenes und nach rechts unbeschränktes Intervall: $\mathbb{R}_0^+ = [0, \infty[$. Das Infimum ist 0 und gleich dem Minimum. Das Supremum ist ∞ . Das Maximum ist nicht definiert.

Betragsfunktion und Dreiecksungleichung

Der Betrag $|x|$ einer reellen Zahl x ist definiert durch

$$|x| = \begin{cases} x & \text{falls } x \geq 0, \\ -x & \text{sonst.} \end{cases} \quad (2.10)$$

Es gilt ferner: $|x| = \sqrt{x^2}$. Die Betragsfunktion ist in Figur 2.1 dargestellt. Das Quadrat des Betrags, dh. $|x|^2$, wird Betragsquadrat genannt. Für reelle Zahlen gilt $|x|^2 = x^2$.

Satz: Für zwei beliebige reelle Zahlen x und y gilt die Dreiecksungleichung:

$$|x \pm y| \leq |x| + |y|. \quad (2.11)$$

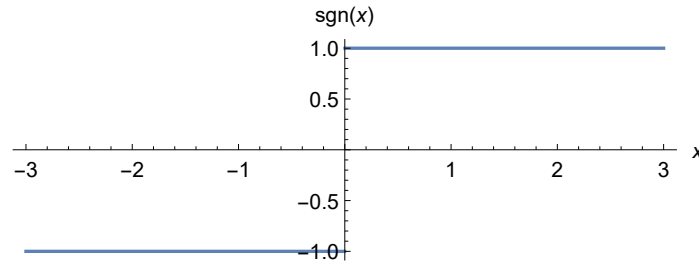


Abbildung 2.2: Die Vorzeichenfunktion (oder Signumfunktion) $f(x) = \operatorname{sgn}(x)$ bildet alle positiven (bzw. negativen) Zahlen auf die Zahl +1 (bzw. -1) ab. Die Zahl 0 wird auf 0 abgebildet.

Beweis: $|x \pm y|^2 = (x \pm y)^2 = x^2 \pm 2xy + y^2 \leq |x|^2 + 2|x||y| + |y|^2 = (|x| + |y|)^2$. Durch Wurzel-Ziehen erhalten wir die Aussage. $\square$

Aufgrund der Dreiecksungleichung gilt $|x \pm y| - |y| \leq |x|$. Wir setzen $a := x \pm y$ und $b := -y$. Das ergibt die umgekehrte Dreiecksungleichung $|a| - |b| \leq |a \pm b|$. Nun benennen wir rück: $x := a$ und $y := b$. Die umgekehrte Dreiecksungleichung und die Dreiecksungleichung kann man dann wie folgt in einer Zeile schreiben:

$$|x| - |y| \leq |x \pm y| \leq |x| + |y| \quad (2.12)$$

Vorzeichenfunktion

Die Vorzeichenfunktion (Figur 2.2) ist definiert durch

$$\operatorname{sgn}(x) = \begin{cases} 1 & \text{falls } x > 0, \\ 0 & \text{falls } x = 0, \\ -1 & \text{falls } x < 0. \end{cases} \quad (2.13)$$

Es gilt für alle $x, y \in \mathbb{R}$: $x = |x| \operatorname{sgn}(x)$, $\operatorname{sgn}(-x) = -\operatorname{sgn}(x)$, $\operatorname{sgn}(x) \operatorname{sgn}(y) = \operatorname{sgn}(xy)$, $\operatorname{sgn}(\operatorname{sgn}(x)) = \operatorname{sgn}(x)$.

Trigonometrische Funktionen

Die wichtigsten trigonometrischen Funktionen können am Einheitskreis (Figur 2.3) im rechtwinkligen Dreieck definiert werden:

$$\sin \varphi = y, \quad (2.14)$$

$$\cos \varphi = x, \quad (2.15)$$

$$\tan \varphi = \frac{\sin \varphi}{\cos \varphi} = \frac{y}{x}. \quad (2.16)$$

Es gelten folgende Beziehungen ($k \in \mathbb{Z}$):

$$\sin(-\varphi) = -\sin \varphi \quad \text{ungerade Funktion,} \quad (2.17)$$

$$\cos(-\varphi) = \cos \varphi \quad \text{gerade Funktion,} \quad (2.18)$$

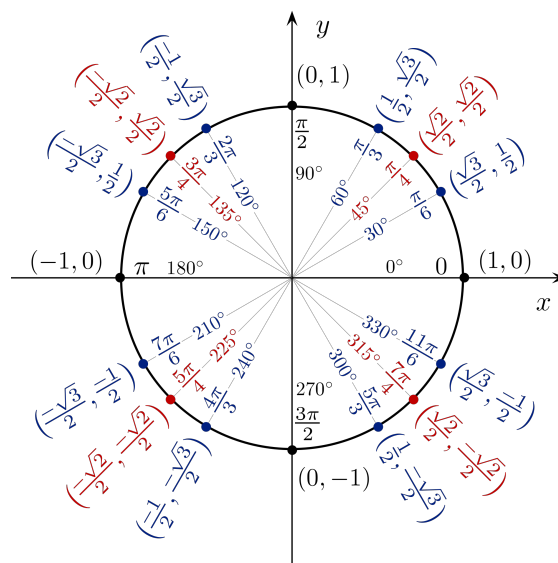


Abbildung 2.3: Einheitskreis. Alle Punkte $P(x, y)$ mit $x = \cos \varphi$ und $y = \sin \varphi$ haben den Abstand 1 vom Ursprung. Bildquelle: Wikipedia.

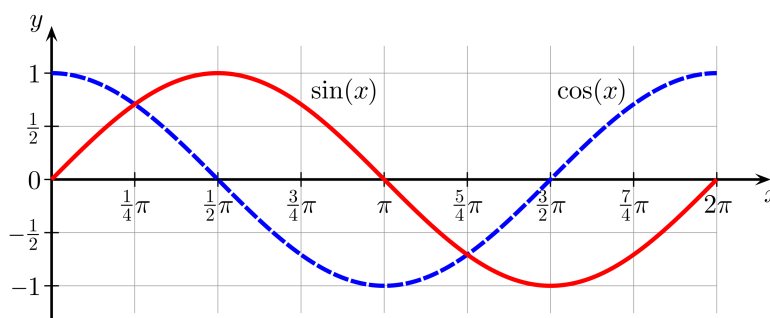


Abbildung 2.4: Sinus (rot) und Cosinus (blau) sind von zentraler Bedeutung für harmonische und periodische Prozesse, insbesondere für Wellen und Schwingungen. Bildquelle: Wikipedia.

$$\sin \varphi = \cos\left(\varphi - \frac{\pi}{2}\right) \quad \text{Phasenverschiebung,} \quad (2.19)$$

$$\sin^2 \varphi + \cos^2 \varphi = 1 \quad \text{trigonometrischer Pythagoras,} \quad (2.20)$$

$$\sin \varphi = \sin(\varphi + 2\pi k) \quad \text{Periodizität,} \quad (2.21)$$

$$\cos \varphi = \cos(\varphi + 2\pi k) \quad \text{Periodizität.} \quad (2.22)$$

Figur 2.4 illustriert die Funktionen Sinus und Cosinus im Intervall $[0, 2\pi]$.

2.2 Komplexe Zahlen

Der Begriff “komplexe Zahlen” stammt von Carl Friedrich Gauß (*Theoria residuorum biquadraticorum*, 1831). Der Ursprung der Theorie der komplexen Zahlen geht auf Gerolamo Cardano (*Ars magna*, 1545) und Rafael Bombelli (*L’Algebra*, 1572) zurück.

Die komplexen Zahlen $\mathbb{C}$ stellen eine Erweiterung der reellen Zahlen dar, die es möglich macht, Gleichungen der Form $x^2 = -1$ zu lösen. Im Gegensatz zu den Erweiterungen von

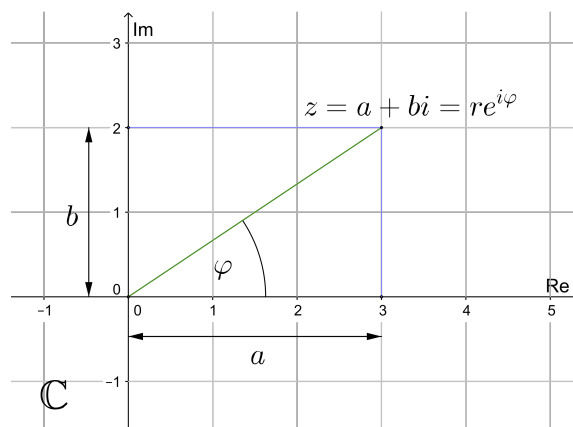


Abbildung 2.5: Die komplexe (Gaußsche) Zahlenebene. Jede komplexe Zahl $z = a + ib$ hat einen Realteil (x -Koordinate) und einen Imaginärteil (y -Koordinate). Die Polardarstellung $z = re^{i\varphi}$ verwendet den Radius (Abstand vom Ursprung) $r = \sqrt{a^2 + b^2}$ und den Winkel $\varphi = \arg(z)$. Bildquelle: Wikipedia.

$\mathbb{N} \rightarrow \mathbb{Z}$ (linksseitig erweitern) und $\mathbb{Z} \rightarrow \mathbb{Q} \rightarrow \mathbb{R}$ (verdichten), ist es nötig, die Zahlengerade zu verlassen und auf eine Zahlenebene zu erweitern.

Komplexe Zahlen haben die Form

$$z = a + ib. \quad (2.23)$$

Den Realteil einer komplexen Zahl z schreibt man als $\operatorname{Re}(z) = a$, den Imaginärteil als $\operatorname{Im}(z) = b$, wobei a und b reelle Zahlen sind. Dargestellt werden können komplexe Zahlen in der Gaußschen (komplexen) Zahlenebene. Der Realteil wird auf der x -Achse aufgetragen, der Imaginärteil auf der y -Achse (Figur 2.5). Die Zahl

$$i := \sqrt{-1} \quad (2.24)$$

wird als imaginäre Einheit bezeichnet. Komplexe Zahlen ohne Imaginärteil ($b = 0$) sind reell. Komplexe Zahlen ohne Realteil ($a = 0$) sind imaginär.

Die komplexen Zahlen $\mathbb{C}$ bilden bezüglich Addition und Multiplikation einen Körper.¹ Zwei komplexe Zahlen $z_1 = a_1 + ib_1$ und $z_2 = a_2 + ib_2$ werden wie folgt addiert bzw. multipliziert:

$$z_1 + z_2 = a_1 + ib_1 + a_2 + ib_2 = (a_1 + a_2) + i(b_1 + b_2), \quad (2.25)$$

$$z_1 z_2 = (a_1 + ib_1)(a_2 + ib_2) = (a_1 a_2 - b_1 b_2) + i(a_1 b_2 + a_2 b_1). \quad (2.26)$$

Die Addition entspricht geometrisch der Vektoraddition, also dem Aneinanderlegen von Vektorpfeilen. Die jeweiligen Komponenten (Koordinaten) werden einfach addiert. Die Multiplikation können wir in dieser Schreibweise nur schwerlich interpretieren. Das holen

¹Kommutativgesetz, Assoziativgesetz und Distributivgesetz gelten. Neutrale Elemente 0 (bezüglich Addition) und 1 (bezüglich Multiplikation). Inverse Elemente $-z$ (bezüglich Addition) und $z^{-1} = \frac{1}{z} = \frac{\bar{z}}{|z|^2}$ (bezüglich Multiplikation).

wir später nach, wenn wir zur Polardarstellung übergehen.

Konjugation

Die zu $z = a + ib$ komplex konjugierte (kurz: konjugierte) Zahl ist

$$\bar{z} = a - ib. \quad (2.27)$$

Die Konjugation entspricht also einer Spiegelung an der x -Achse.

Das Produkt einer komplexen Zahl mit ihrer konjugierten Zahl ist

$$z\bar{z} = (a + ib)(a - ib) = a^2 - i^2 b^2 = a^2 + b^2 = |z|^2, \quad (2.28)$$

also das Betragsquadrat von z . Damit kann für jede Zahl $z \neq 0$ ihr Inverses wie folgt berechnet werden: $z^{-1} = \bar{z}/|z|^2$. Ferner gilt (Beweis in den Übungen):

$$\overline{z_1 + z_2} = \bar{z}_1 + \bar{z}_2, \quad (2.29)$$

$$\overline{z_1 z_2} = \bar{z}_1 \bar{z}_2. \quad (2.30)$$

Daraus folgt:

$$|z_1 z_2|^2 = z_1 z_2 \overline{z_1 z_2} = z_1 \bar{z}_1 z_2 \bar{z}_2 = |z_1|^2 |z_2|^2 \Rightarrow |z_1 z_2| = |z_1| |z_2|. \quad (2.31)$$

Die Summe einer komplexen Zahl mit ihrer Konjugierten ist eine reelle Zahl:

$$z + \bar{z} = a + ib + a - ib = 2a = 2\operatorname{Re}(z). \quad (2.32)$$

Polardarstellung

Eine komplexe Zahl $z = a + ib$ in kartesischer (auch genannt: algebraischer) Darstellung lässt sich in Polarform umschreiben, indem wir Polarkoordinaten mit Radius $r \in \mathbb{R}_0^+ = \{x \in \mathbb{R} | x \geq 0\} = [0, \infty[$ und Winkel $\varphi \in [0, 2\pi[$ einführen:

$$a = r \cos \varphi, \quad (2.33)$$

$$b = r \sin \varphi. \quad (2.34)$$

Damit erhalten wir:

$$z = r(\cos \varphi + i \sin \varphi) = r e^{i\varphi}. \quad (2.35)$$

Die Eulersche Formel $e^{i\varphi} = \cos \varphi + i \sin \varphi$ werden wir später beweisen; dafür fehlt uns jetzt noch das Rüstzeug. Der Ausdruck $e^{i\varphi}$ beschreibt den komplexen Einheitskreis, also einen Kreis mit Radius 1 in der Gaußschen Ebene.

In Polardarstellung nimmt die Multiplikation zweier komplexer Zahlen $z_1 = r_1 e^{i\varphi_1}$

und $z_2 = r_2 e^{i\varphi_2}$ eine besonders einfache und anschauliche Form an:

$$z_1 z_2 = r_1 e^{i\varphi_1} r_2 e^{i\varphi_2} = r_1 r_2 e^{i(\varphi_1 + \varphi_2)}. \quad (2.36)$$

Bei der Multiplikation werden die Radien multipliziert und die Winkel addiert.

Wurzeln

Gegeben eine komplexe Zahl $c = |c| e^{i\phi}$. Die Lösungen der Gleichung

$$z^n = c \quad (2.37)$$

bezeichnet man als die n -ten Wurzeln von c :

$$z_k = \sqrt[n]{|c|} \exp\left(\frac{i\phi}{n} + k \frac{2\pi i}{n}\right), \quad k = 0, 1, 2, \dots, n-1. \quad (2.38)$$

Beispiel: Berechne die dritten Wurzeln der Zahl $1 = 1 e^{i0}$, dh. finde die Lösungen der Gleichung $z^3 = 1$. *Lösung:* $z_0 = 1$, $z_1 = \exp(\frac{2\pi i}{3})$, $z_2 = \exp(\frac{4\pi i}{3})$.

Ordnung und Vollständigkeit

Satz: Es gibt auf den komplexen Zahlen keine Ordnung (und daher auch keine Ordnungsvollständigkeit).

Beweis durch Widerspruch: Nehmen wir an, die Ordnungsaxiome wären erfüllt, insbesondere die Totalität und die Verträglichkeit mit Multiplikation. Betrachten wir die Zahl i . Fall 1: $i \geq 0$. Wir multiplizieren beide Seiten mit i und erhalten $i^2 \geq 0$, was falsch ist. Fall 2: $i \leq 0$. Daraus folgt $-i \geq 0$. Wir multiplizieren beide Seiten mit $-i$ und erhalten erneut die falsche Aussage $i^2 \geq 0$. Also sind die Ordnungsaxiome verletzt. Die komplexen Zahlen sind mithin nicht geordnet. $\square$

Die komplexen Zahlen sind allerdings – im Unterschied zu den reellen Zahlen – algebraisch vollständig. Es gilt nämlich der Gauß-d'Alembertsche *Fundamentalsatz der Algebra*: Jedes nicht konstante² Polynom besitzt im Bereich der komplexen Zahlen mindestens eine Nullstelle.

Beispiel: Das Polynom $P(x) = x^2 + 1$ hat keine Nullstelle in den reellen Zahlen. Aber zwei Nullstellen in den komplexen Zahlen, nämlich $x = i$ und $x = -i$.

²Konstante Polynome haben die Form $P(x) = c$ und haben in Ermangelung einer Abhängigkeit von x keine Nullstellen, außer das Nullpolynom $P(x) = 0$.

3. Folgen und Reihen

3.1 Folgen

Die Theorie der Grenzwerte unendlicher Folgen ist von zentraler Bedeutung für die Analysis, denn auf ihr fußt der Grenzwert von Funktionen, die Definition der Ableitung sowie der Riemannsche Integralbegriff.

Definition: Eine *Folge* ist eine Auflistung von fortlaufenden nummerierten Objekten, insbesondere von Zahlen.

Die einzelnen Objekte nennt man die Komponenten oder die Glieder der Folge. Im Gegensatz zu den Elementen einer Menge ist bei Folgen die Reihenfolge der Komponenten relevant und diese können sich auch wiederholen.

Beispiele: Die Folge $(1, 2, *, \square, 1)$ ist eine endliche Folge der Länge 5. Die Folge $(2, 3, 5, 7, 11, 13, 17, \dots)$ ist die unendliche Zahlenfolge aller Primzahlen.

Allgemein schreiben wir für eine endliche bzw. unendliche Folge

$$(a_n)_{n=1, \dots, N} = (a_1, a_2, \dots, a_N), \quad (3.1)$$

$$(a_n)_{n \in \mathbb{N}} = (a_1, a_2, \dots). \quad (3.2)$$

Eine (unendliche) Folge ist daher eine Abbildung der natürlichen Zahlen auf die Komponenten der Folge,

$$a : \mathbb{N} \rightarrow Y, n \mapsto a_n \quad (3.3)$$

wobei Y die Ziel- bzw. Bildmenge ist.

Bekannte (unendliche) Folgen mit $Y = \mathbb{R}$ bzw. $Y = \mathbb{N}$ sind:

- Die arithmetische Folge: $a_n = a_1 + (n-1)d$, dh. $(a_n)_{n \in \mathbb{N}} = (a_1, a_1 + d, a_1 + 2d, a_1 + 3d, \dots)$. Die Differenz zwischen zwei aufeinanderfolgenden Folgengliedern ist konstant: $a_{n+1} - a_n = d$. Beispiel: Mit $a_1 = 1$ und $d = 1$ erhält man $(a_n)_{n \in \mathbb{N}} =$

$(1, 2, 3, 4, \dots)$.

- Die harmonische Folge (Kehrwerte der natürlichen Zahlen): $a_n = \frac{1}{n}$, dh. $(a_n)_{n \in \mathbb{N}} = (1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots)$.
- Geometrische Folgen (exponentielle/s Abschwächung/Wachstum): $a_n = a_1 q^{n-1}$, dh. $(a_n)_{n \in \mathbb{N}} = (a_1, a_1 q, a_1 q^2, a_1 q^3, \dots)$. Der Quotient zwischen zwei aufeinanderfolgenden Folgengliedern ist konstant: $\frac{a_{n+1}}{a_n} = q$. Beispiele: Mit $a_1 = 1$ und $q = 2$ erhält man $(a_n)_{n \in \mathbb{N}} = (1, 2, 4, 8, 16, \dots)$. Für $a_1 = 1$ und $q = \frac{1}{2}$ ist $(a_n)_{n \in \mathbb{N}} = (1, \frac{1}{2}, \frac{1}{4}, \frac{1}{8}, \frac{1}{16}, \dots)$.
- Die Fibonacci-Folge (Summe der beiden Vorgänger): $a_n = a_{n-2} + a_{n-1}$ mit $a_1 = 1, a_2 = 1$, dh. $(a_n)_{n \in \mathbb{N}} = (1, 1, 2, 3, 5, 8, 13, \dots)$.
- Verzinsung mit Zinssatz r nach n Jahren: $a_n = (1 + z)^n$. Beispiel: Mit Zinssatz $r = 0.1$ ergibt sich $(a_n)_{n \in \mathbb{N}} = (1.10, 1.21, 1.33, 1.46, 1.61, 1.77, \dots)$.

Konvergenz und Divergenz

Definition: Die Folge $(a_n)_{n \in \mathbb{N}}$ konvergiert gegen a für $n \rightarrow \infty$ falls

$$\forall \varepsilon > 0 \exists N_\varepsilon \in \mathbb{N} \forall n \geq N_\varepsilon : |a_n - a| < \varepsilon, \quad (3.4)$$

also falls für jede beliebig kleine (positive) Fehlerschranke ε ab einem bestimmten Folgenglied mit ausreichend großem Index N_ε alle Folgenglieder um weniger als ε von a abweichen.

Salopp formuliert: Späte Folgenglieder liegen beliebig dicht bei a . Geschrieben wird dies als

$$a = \lim_{n \rightarrow \infty} a_n, \quad (3.5)$$

wobei a der *Grenzwert* oder *Limes* der Folge genannt wird. Alternativ schreibt man auch: $a_n \rightarrow a$ für $n \rightarrow \infty$ (gesprochen als: “ a_n strebt gegen a für n gegen unendlich”).

Eine Folge heißt *konvergent*, wenn sie einen Limes besitzt. Ansonsten nennen wir die Folge *divergent*.

Beispiel: Die harmonische Folge $a_n = \frac{1}{n}$ konvergiert gegen den Grenzwert $a = 0$. *Beweis:* Es gilt zunächst

$$|a_n - a| = \left| \frac{1}{n} - 0 \right| = \frac{1}{n}. \quad (3.6)$$

Für jedes $\varepsilon > 0$ kann man eine natürliche Zahl wählen, die größer als der Kehrwert ist, dh. $N_\varepsilon > \frac{1}{\varepsilon}$, z.B. $N_\varepsilon := \lfloor \frac{1}{\varepsilon} \rfloor + 1$. Dann gilt für alle $n \geq N_\varepsilon$:

$$|a_n - a| = \frac{1}{n} \leq \frac{1}{N_\varepsilon} < \varepsilon. \quad (3.7)$$

Also ist $\lim_{n \rightarrow \infty} \frac{1}{n} = 0$. $\square$ Figur 3.1 zeigt die harmonische Folge und eine ε -Umgebung.

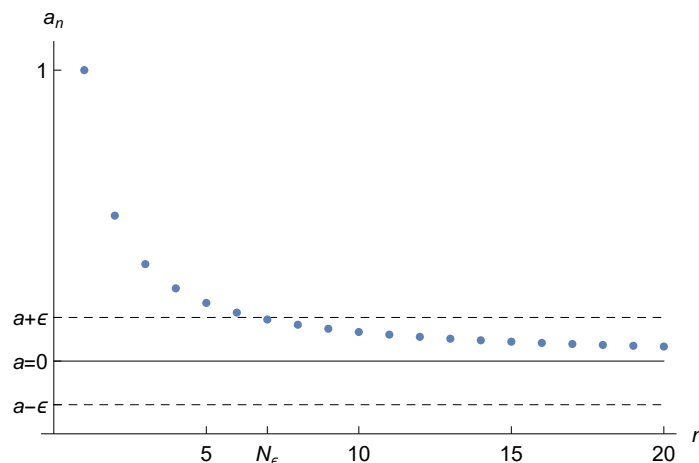


Abbildung 3.1: Die harmonische Zahlenfolge $a_n = \frac{1}{n}$. Für jedes $\varepsilon > 0$ kann man ein N_ε finden, sodass alle Folgenglieder mit Index $n \geq N_\varepsilon$ im Intervall $]a - \varepsilon, a + \varepsilon[$ liegen, also $|a_n - a| < \varepsilon$ gilt. In der Abbildung ist $\varepsilon = 0.15$ und der kleinstmögliche kritische Index ist $N_\varepsilon = 7$.

Beispiel: Die Folge $a_n = 1 + \frac{1}{\sqrt{n}}$ konvergiert gegen den Grenzwert $a = 1$. *Beweis:*

$$|a_n - a| = \left| 1 + \frac{1}{\sqrt{n}} - 1 \right| = \frac{1}{\sqrt{n}}. \quad (3.8)$$

Aufgrund der Wurzel im Nenner benötigen wir nun für jedes $\varepsilon > 0$ einen anderen (späteren) kritischen Index als im vorigen Beispiel, nämlich $N_\varepsilon > \frac{1}{\varepsilon^2}$, z.B. $N_\varepsilon := \lfloor \frac{1}{\varepsilon^2} \rfloor + 1$. Dann gilt für alle $n \geq N_\varepsilon$:

$$|a_n - a| = \frac{1}{\sqrt{n}} \leq \frac{1}{\sqrt{N_\varepsilon}} < \varepsilon. \quad \square \quad (3.9)$$

Beispiel: Sei $|q| < 1$. Dann gilt $\lim_{n \rightarrow \infty} q^n = 0$. *Beweis:* In den Übungen.

Beispiel: Die Folge $a_n = (-1)^n$ ist divergent. *Beweis:* Für jeden potentiellen Grenzwert a gilt aufgrund der Dreiecksungleichung (2.11) für beliebig große n :

$$|a_n - a| + |a_{n+1} - a| \geq |(a_n - a) - (a_{n+1} - a)| = |a_n - a_{n+1}| = 2. \quad (3.10)$$

Wenn a ein Grenzwert wäre, dann müsste die linke Seite ab einem kritischen Index stets $< 2\varepsilon$ sein, andererseits aber auch ≥ 2 . Für $\varepsilon \leq 1$ ist dies unmöglich. Also hat die Folge keinen Grenzwert und ist mithin divergent. $\square$

Eindeutigkeit des Grenzwerts

Satz: Für konvergente Folgen ist der Grenzwert a eindeutig bestimmt, dh es gibt genau einen Grenzwert.

Beweis durch Widerspruch: Hätte die Folge einen zweiten von a verschiedenen

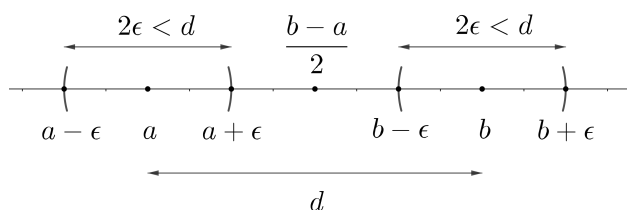


Abbildung 3.2: Hätte eine Folge mit Grenzwert a einen zweiten von a verschiedenen Grenzwert b , so könnte man disjunkte ε -Umgebungen definieren. Ab einem bestimmten Index müssten aber alle Folgenglieder in beiden Intervallen zugleich liegen, was nicht möglich ist. Daher ist der Grenzwert eindeutig und es muss $a = b$ gelten. Bildquelle: Wikipedia.

Grenzwert b , dann können wir den Abstand $d := |a - b| > 0$ definieren. Nun wählen wir $\varepsilon < d/2$ und definieren die beiden ε -Umgebungen zu den beiden Grenzwerten, also die beiden Intervalle $]a - \varepsilon, a + \varepsilon[$ und $]b - \varepsilon, b + \varepsilon[$. Diese Intervalle überlappen nicht (sind disjunkt), haben also keinen Punkt gemeinsam. Ab einem bestimmten Index müssen aber alle Folgenglieder in der ε -Umgebung des Grenzwertes liegen. Dies ist aber nicht für beide Grenzwerte a und b möglich. Dieser Widerspruch lässt sich nur dadurch beheben, dass a und b keinen Abstand haben dürfen, also $a = b$ gelten muss. Also ist der Grenzwert eindeutig. (Veranschaulichung in Figur 3.2.) $\square$

Rechenregeln für Grenzwerte

Für zwei konvergente Folgen $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ mit Grenzwerten $a = \lim_{n \rightarrow \infty} a_n$ bzw. $b = \lim_{n \rightarrow \infty} b_n$ gelten die folgenden Regeln:

- **Summensatz:** $\lim_{n \rightarrow \infty} (a_n + b_n) = a + b$
- **Produktsatz:** $\lim_{n \rightarrow \infty} (a_n b_n) = a b$
- **Quotientensatz:** $\lim_{n \rightarrow \infty} \frac{a_n}{b_n} = \frac{a}{b}$ falls $b \neq 0$ und $\forall n : b_n \neq 0$

Der Faktorsatz $\lim_{n \rightarrow \infty} (\lambda a_n) = \lambda a$ folgt aus dem Produktsatz, wenn man die konstante Folge $b_n = \lambda$ wählt.

Beispiel: Berechne den Grenzwert der Folge $a_n = \frac{3n^2 + 5n + 7}{2n^2 - 4n + 6}$. *Lösung:* Mittels Produkt-, Quotienten- und Summensatz:

$$\begin{aligned}
 \lim_{n \rightarrow \infty} \frac{3n^2 + 5n + 7}{2n^2 - 4n + 6} &= \lim_{n \rightarrow \infty} \frac{3 + \frac{5}{n} + \frac{7}{n^2}}{2 - \frac{4}{n} + \frac{6}{n^2}} = \lim_{n \rightarrow \infty} \left(3 + \frac{5}{n} + \frac{7}{n^2}\right) \lim_{n \rightarrow \infty} \frac{1}{2 - \frac{4}{n} + \frac{6}{n^2}} \\
 &= \lim_{n \rightarrow \infty} \left(3 + \frac{5}{n} + \frac{7}{n^2}\right) \frac{1}{\lim_{n \rightarrow \infty} 2 - \frac{4}{n} + \frac{6}{n^2}} \\
 &= \left(\lim_{n \rightarrow \infty} 3 + \lim_{n \rightarrow \infty} \frac{5}{n} + \lim_{n \rightarrow \infty} \frac{7}{n^2}\right) \frac{1}{\lim_{n \rightarrow \infty} 2 - \lim_{n \rightarrow \infty} \frac{4}{n} + \lim_{n \rightarrow \infty} \frac{6}{n^2}} \\
 &= (3 + 0 + 0) \frac{1}{2 - 0 + 0} = \frac{3}{2}. \tag{3.11}
 \end{aligned}$$

Ganz allgemein sind bei einem Polynombruch immer die führenden Potenzen entscheidend.

Uneigentliche Grenzwerte

Strebt eine divergente Folge $(a_n)_{n \in \mathbb{N}}$ für große n zu unbegrenzt großen Werten, also falls gilt

$$\forall A \in \mathbb{R} \exists n_A \in \mathbb{N} \forall n \geq n_A : a_n \geq A, \quad (3.12)$$

dann hat die Folge den *uneigentlichen* Grenzwert ∞ und man schreibt

$$\lim_{n \rightarrow \infty} a_n = \infty. \quad (3.13)$$

Analog schreiben wir $\lim_{n \rightarrow \infty} a_n = -\infty$, falls die Folge für große n zu unbegrenzt kleinen Werten strebt.

Wenn eine Folge einen uneigentlichen Grenzwert hat, dann nennt man diese divergente Folge auch “uneigentlich konvergent” oder “bestimmt divergent”.

Beispiel: Die Fibonacci-Folge $(a_n)_{n \in \mathbb{N}} = (1, 1, 2, 3, 5, 8, \dots)$ ist divergent mit dem uneigentlichen Grenzwert $\lim_{n \rightarrow \infty} a_n = \infty$.

Beispiel: Wir betrachten die Folge $a_n = q^n$. Dann gilt:

$$\lim_{n \rightarrow \infty} a_n = \begin{cases} 0 & \text{falls } |q| < 1 \text{ (konvergent),} \\ 1 & \text{falls } q = 1 \text{ (konvergent),} \\ \infty & \text{falls } q > 1 \text{ (divergent).} \end{cases} \quad (3.14)$$

Für $q \leq -1$ ist die Folge divergent, und es gibt auch keinen uneigentlichen Grenzwert.

Monotonie und Beschränktheit

Eine Folge $(a_n)_{n \in \mathbb{N}}$ nennen wir

- *monoton wachsend*, wenn gilt: $\forall n \in \mathbb{N} : a_n \leq a_{n+1}$,
- *monoton fallend*, wenn gilt: $\forall n \in \mathbb{N} : a_n \geq a_{n+1}$,
- *nach oben beschränkt*, wenn gilt: $\exists T \in \mathbb{R} \forall n : a_n \leq T$,
- *nach unten beschränkt*, wenn gilt: $\exists T \in \mathbb{R} \forall n : a_n \geq T$.

Ein Folge heißt *streng* monoton wachsend bzw. fallend, wenn in den ersten beiden Punkten die strikte Ungleichung $a_n < a_{n+1}$ bzw. $a_n > a_{n+1}$ gilt.

Die für eine nach oben (bzw. unten) beschränkte Folge kleinst (bzw. größt) mögliche Schranke T bezeichnet man als Supremum (bzw. Infimum) der Folge und schreibt $\sup_{n \in \mathbb{N}} a_n$ (bzw. $\inf_{n \in \mathbb{N}} a_n$).

Eine Folge heißt beschränkt, wenn sie sowohl nach oben als auch nach unten beschränkt ist.

Satz: Jede konvergente Folge ist beschränkt. Und im Umkehrschluss: Jede unbeschränkte Folge ist divergent.^a

Beweis: Konvergenz bedeutet, dass $|a_n - a| < \varepsilon$ für hinreichend späte Folgenglieder

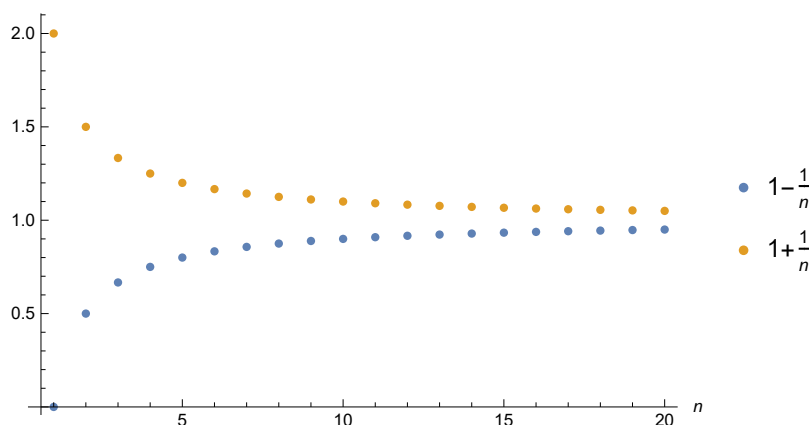


Abbildung 3.3: Die Folge $a_n = 1 - \frac{1}{n}$ (blau) ist monoton wachsend und beschränkt und hat den Grenzwert 1. Die Folge $b_n = 1 + \frac{1}{n}$ (orange) ist monoton fallend und beschränkt und hat auch den Grenzwert 1.

$n \geq N_\varepsilon$. Dh. für $\varepsilon = 1$ gilt, dass alle a_n mit $n \geq N_\varepsilon =: M$ im Intervall $]a - 1, a + 1[$ liegen. Alle davor liegenden Folgenglieder mit $n \leq M - 1$ können durch das Maximum bzw. Minimum abgeschätzt werden. Insgesamt gilt für alle $n \in \mathbb{N}$: $\min(a_1, a_2, \dots, a_{M-1}, a - 1) \leq a_n \leq \max(a_1, a_2, \dots, a_{M-1}, a + 1)$. Also ist die Folge beschränkt. $\square^b$

^aDaraus folgt nicht, dass jede beschränkte Folge konvergent ist. Und es folgt nicht, dass jede divergente Folge unbeschränkt ist.

^bMan könnte versucht sein, zu behaupten, dass die Folge $(a_n)_{n \in \mathbb{N}}$ mit $a_n = \frac{1}{n-1}$ ein Beispiel für eine unbeschränkte (weil $a_1 = \frac{1}{0}$) und konvergente Folge (weil $\lim_{n \rightarrow \infty} a_n = 0$) ist, im Widerspruch zum soeben bewiesenen Satz. Aber $\frac{1}{0}$ ist keine reelle Zahl; das Folgenglied a_1 ist unzulässig definiert. Also ist dies keine erlaubte Folge.

Beispiele:

- Die Folge $(a_n)_{n \in \mathbb{N}}$ mit $a_n = 1 - \frac{1}{n}$ ist monoton wachsend und beschränkt mit $\sup_{n \in \mathbb{N}} a_n = 1$ und $\inf_{n \in \mathbb{N}} a_n = 0$. Der Grenzwert der Folge ist $\lim_{n \rightarrow \infty} a_n = \lim_{n \rightarrow \infty} (1 - \frac{1}{n}) = 1 - 0 = 1$.
- Die Folge $(b_n)_{n \in \mathbb{N}}$ mit $b_n = 1 + \frac{1}{n}$ ist monoton fallend und beschränkt mit $\sup_{n \in \mathbb{N}} b_n = 2$ und $\inf_{n \in \mathbb{N}} b_n = 1$. Der Grenzwert der Folge ist $\lim_{n \rightarrow \infty} b_n = \lim_{n \rightarrow \infty} (1 + \frac{1}{n}) = 1 + 0 = 1$.

Figur 3.3 veranschaulicht diese beiden Beispiele.

Monotoniekriterium

Satz (Monotoniekriterium): Für jede monoton wachsende Folge $(a_n)_{n \in \mathbb{N}}$ reeller Zahlen gilt: Die Folge konvergiert genau dann, wenn sie nach oben beschränkt ist. Der Grenzwert ist dann $\lim_{n \rightarrow \infty} a_n = \sup_{n \in \mathbb{N}} a_n$.

Beweis: Von links nach rechts (nach oben beschränkt $\Rightarrow$ konvergent) direkt: Aus der Beschränktheit folgt, dass die Menge aller Folgenglieder $A = \{a_n | n \in \mathbb{N}\}$ ein Supremum

(kleinste obere Schranke) $\sup A = a$ besitzt. Wir wählen ein $\varepsilon > 0$. Da die Folge monoton wächst und es keine kleinere obere Schranke als a gibt, gilt ab einem gewissen Index N für alle $a_n \geq N$ die Ungleichung $a - \varepsilon < a_n \leq a$. Also ist $|a_n - a| = a - a_n < \varepsilon$, dh. die Folge konvergiert und hat als Grenzwert das Supremum a .

Von rechts nach links (nach oben beschränkt $\Leftrightarrow$ konvergent)^a indirekt durch Widerspruch: Wir nehmen an, dass trotz Konvergenz, dh. trotz $\lim_{n \rightarrow \infty} a_n = a$ die Folge nach oben unbeschränkt ist, es also ein Folgenglied a_N gibt, sodass $a_N > a$. Da die Folge monoton wächst, gilt dann für alle $n \geq N$, dass $a_n \geq a_N > a$ bzw. $a_n - a \geq a_N - a > 0$. Nun können wir $\varepsilon = a_N - a$ wählen. Dann gilt für alle $n \geq N$, dass $a_n - a = |a_n - a| \geq \varepsilon$. Dies widerspricht der angenommenen Konvergenz. Also muss die Annahme der Unbeschränktheit falsch gewesen sein. Mithin ist die Folge nach oben beschränkt. $\square$

^aDiese Richtung wurde schon im vorhergehenden Satz direkt bewiesen, der besagt, dass jede konvergente Folge beschränkt ist. Hier wird nun eine weitere Beweismethode gezeigt.

Analog dazu konvergiert eine monoton fallende Folge genau dann, wenn sie nach unten beschränkt ist. Ihr Grenzwert ist dann das Infimum der Menge aller Folgenglieder.

Beispiel: Die Folge

$$a_n = \left(1 + \frac{1}{n}\right)^n \quad (3.15)$$

ist monoton wachsend und beschränkt. *Beweis:* In den Übungen. Die Folge konvergiert gegen die Eulersche Zahl e : $\lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n = e$.

Teilfolgen und Häufungspunkte

Definition: Wir nennen $(a_{n_k})_{k \in \mathbb{N}} = (a_{n_1}, a_{n_2}, a_{n_3}, \dots)$ eine *Teilfolge* der Folge $(a_n)_{n \in \mathbb{N}}$, wenn $n_1 < n_2 < n_3 < \dots$ gilt. Falls eine Teilfolge mit Grenzwert a' existiert, dann nennen wir a' einen *Häufungspunkt* der Folge $(a_n)_{n \in \mathbb{N}}$.

Beispiel: Die (divergente) Folge $a_n = (-1)^n = (-1, +1, -1, +1, \dots)$ hat zwei konvergente Teilfolgen $a'_n = (-1, -1, -1, \dots)$ und $a''_n = (+1, +1, +1, \dots)$. Die beiden Häufungspunkte von a_n sind daher $a' = -1$ und $a'' = +1$.

Den größten bzw. kleinsten Häufungspunkt einer Folge reeller Zahlen bezeichnet man als *Limes superior* ($\limsup$) bzw. *Limes inferior* ($\liminf$):

$$\limsup_{n \rightarrow \infty} a_n := \inf_{n \in \mathbb{N}} \sup_{m \geq n} a_m, \quad (3.16)$$

$$\liminf_{n \rightarrow \infty} a_n := \sup_{n \in \mathbb{N}} \inf_{m \geq n} a_m. \quad (3.17)$$

Sie stellen einen Ersatz für den Grenzwert dar, wenn dieser nicht existiert (siehe Figur 3.4). Falls der Grenzwert existiert, dann gibt es nur einen Häufungspunkt und es gilt $\lim_{n \rightarrow \infty} a_n = \limsup_{n \rightarrow \infty} a_n = \liminf_{n \rightarrow \infty} a_n$.

Falls eine Folge nicht nach oben (bzw. unten) beschränkt ist, dann gibt es keinen

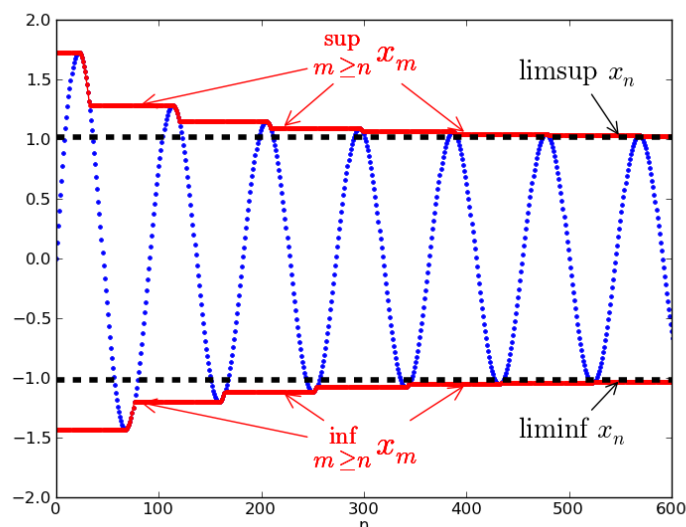


Abbildung 3.4: Die Folge x_n (blau) ist beschränkt und divergent. Trotz Divergenz besitzt sie (wegen ihrer Beschränktheit) einen Limes superior und einen Limes inferior. Für jedes n stellt die rote Kurve das Supremum (oben) bzw. Infimum (unten) für alle $m \geq n$ dar. Das Supremum sinkt und die größte untere Schranke des Supremums ist $\inf_{n \in \mathbb{N}} \sup_{m \geq n} x_m =: \limsup_{n \rightarrow \infty} x_n$, der Limes superior. Ganz analog dazu steigt das Infimum und die kleinste obere Schranke des Infimums ist $\sup_{n \in \mathbb{N}} \inf_{m \geq n} x_m =: \liminf_{n \rightarrow \infty} x_n$, der Limes inferior. Bildquelle: Wikipedia.

größten (bzw. kleinsten) Häufungspunkt. Den Limes superior (bzw. inferior) definiert man dann als ∞ (bzw. $-\infty$).

Satz von Bolzano-Weierstraß (ohne Beweis): Jede beschränkte Folge $(a_n)_{n \in \mathbb{N}}$ enthält (mindestens) eine konvergente Teilfolge, besitzt also auch einen Häufungspunkt.

Beispiel: Die Folgen $a_n = \frac{1}{2}(1 - \frac{2}{n})$ und $b_n = \sin \frac{n\pi}{20}$ sind beide beschränkt (Figur 3.5). Die Folge a_n ist monoton wachsend und konvergiert nach dem Monotoniekriterium gegen ihr Supremum $\frac{1}{2}$. Dies ist auch der einzige Häufungspunkt der Folge und damit sowohl der Limes superior als auch der Limes inferior. Die Folge b_n ist divergent, hat also keinen Grenzwert. Aber sie besitzt viele Häufungspunkte (genau gesagt: 21 Stück). Ein solcher Häufungspunkt ist beispielsweise 0, da die Teilfolge $(\sin \frac{20\pi}{20}, \sin \frac{40\pi}{20}, \sin \frac{60\pi}{20}, \dots) = (0, 0, 0, \dots)$ gegen 0 konvergiert. Der größte Häufungspunkt ist der Limes superior mit Wert 1. Der kleinste Häufungspunkt ist der Limes inferior mit dem Wert -1 .

Cauchy-Kriterium

Das Cauchy-Kriterium erlaubt es, auch ohne Kenntnis des Grenzwerts zu entscheiden, ob eine Folge konvergent oder divergent ist.

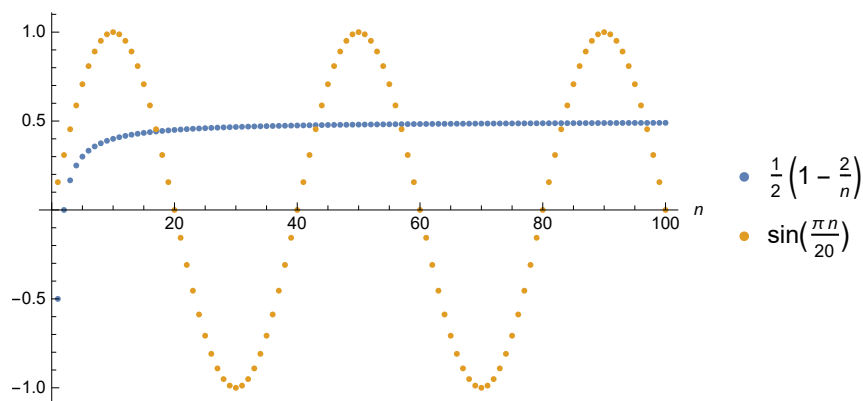


Abbildung 3.5: Die Folgen $\frac{1}{2}\left(1 - \frac{2}{n}\right)$ (blau) und $\sin\frac{n\pi}{20}$ (orange) sind beide beschränkt. Nach dem Satz von Bolzano-Weierstraß haben sie daher mindestens einen Häufungspunkt. Während die blaue Folge konvergent ist und daher nur einen Häufungspunkt hat, hat die orange Folge viele Häufungspunkte.

Definition: Eine Folge $(a_n)_{n \in \mathbb{N}}$ heißt *Cauchy-Folge*, wenn gilt

$$\forall \varepsilon > 0 \exists N_\varepsilon \in \mathbb{N} \forall n, m \geq N_\varepsilon : |a_n - a_m| < \varepsilon, \quad (3.18)$$

also falls für jede Fehlerschranke ε ab einem bestimmten Folgenglied mit ausreichend großem Index N_ε alle Folgenglieder um weniger als ε voneinander abweichen.

Vereinfacht formuliert: Bei einer Cauchy-Folge liegen alle späten Folgenglieder beliebig dicht beieinander.

Satz: Eine Folge $(a_n)_{n \in \mathbb{N}}$ ist genau dann konvergent, wenn sie eine Cauchy-Folge ist.

Beweis: Von links nach rechts (Konvergenz $\Rightarrow$ Cauchy-Folge): Konvergenz bedingt die Existenz eines Grenzwerts $a = \lim_{n \rightarrow \infty} a_n$ sowie, dass es zu jedem ε ein N_ε gibt, sodass $\forall n \geq N_\varepsilon : |a_n - a| < \frac{\varepsilon}{2}$.^a Daraus folgt (mithilfe der Dreiecksungleichung) für alle $n, m \geq N_\varepsilon$:

$$|a_n - a_m| = |(a_n - a) - (a_m - a)| \leq |a_n - a| + |a_m - a| < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon. \quad (3.19)$$

Von rechts nach links (Konvergenz $\Leftarrow$ Cauchy-Folge): Zunächst bemerken wir, dass jede Cauchy-Folge beschränkt ist. Zu $\varepsilon = 1$ existiert ein $N_\varepsilon =: M$, sodass für $n, m \geq M$ stets $|a_n - a_m| < \varepsilon = 1$ gilt. Wegen der Dreiecksungleichung gilt $|a_n| - |a_m| \leq |a_n - a_m| < 1$, dh. $|a_n| < 1 + |a_m|$. Ab Index M , dh. für $n \geq M$ sind also alle Folgenglieder a_n beschränkt mit $|a_n| < 1 + |a_M|$. Und für $n \leq M - 1$ kann man einfach das Maximum der Beträge nehmen: Daraus folgt für alle $n \in \mathbb{N}$, dass $|a_n| \leq \max(|a_1|, |a_2|, \dots, |a_{M-1}|, |a_M| + 1)$. Also ist $(a_n)_{n \in \mathbb{N}}$ beschränkt. Nach Bolzano-Weierstraß gibt es wegen der Beschränktheit eine Teilfolge $(a'_n)_{n \in \mathbb{N}}$ mit einem Häufungspunkt, den wir a nennen. Nun müssen wir nur noch zeigen, dass sogar die Gesamtfolge $(a_n)_{n \in \mathbb{N}}$ gegen a strebt. Wir wählen $\varepsilon > 0$ und ein N_ε , sodass (weil Cauchy-Folge) gilt: $|a_n - a_m| < \frac{\varepsilon}{2}$ für alle $n, m \geq N_\varepsilon$. Nun

wählen wir ein Folgenglied a_N mit $N \geq N_\varepsilon$ und $|a_N - a| < \frac{\varepsilon}{2}$. Ein solches Folgenglied existiert, weil die Teilfolge gegen a konvergiert. Für alle $n \geq N_\varepsilon$ ist dann (mithilfe der Dreiecksungleichung):

$$|a_n - a| = |(a_n - a_N) + (a_N - a)| \leq |a_n - a_N| + |a_N - a| < \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon. \quad (3.20)$$

Also konvergiert a_n gegen a . $\square$

^aDer Faktor $\frac{1}{2}$ vereinfacht die Beweisführung.

3.2 Reihen

Wir betrachten die Folge $(a_k)_{k \in \mathbb{N}}$. Wir können die Folgenglieder bis zum n -ten Glied aufsummieren, also Partialsummen (Teilsommen) bilden:

$$\begin{aligned} S_1 &:= a_1, \\ S_2 &:= a_1 + a_2, \\ S_3 &:= a_1 + a_2 + a_3, \\ &\vdots \\ S_n &:= a_1 + a_2 + a_3 + \dots + a_n = \sum_{k=1}^n a_k. \end{aligned} \quad (3.21)$$

Diese Partialsummen bilden selber wieder eine Folge

$$(S_n)_{n \in \mathbb{N}} = (S_1, S_2, S_3, \dots) = (a_1, a_1 + a_2, a_1 + a_2 + a_3, \dots). \quad (3.22)$$

Definition: Eine *Reihe* $(S_n)_{n \in \mathbb{N}}$ (oder Summenfolge oder unendliche Summe oder unendliche Reihe) ist die (unendliche) Folge der n -ten Partialsummen $S_n = \sum_{k=1}^n a_k$ einer Folge $(a_k)_{k \in \mathbb{N}}$. Für die Reihe schreibt man symbolisch

$$\sum_{k=1}^{\infty} a_k. \quad (3.23)$$

Konvergenz und Divergenz

Definition: Die Reihe $\sum_{k=1}^{\infty} a_k$ ist konvergent, falls

$$\lim_{n \rightarrow \infty} S_n = \lim_{n \rightarrow \infty} \sum_{k=1}^n a_k = \sum_{k=1}^{\infty} a_k \quad (3.24)$$

existiert. Dann existiert der *Wert* (auch genannt der Grenzwert oder die Summe) $S = \lim_{n \rightarrow \infty} S_n$ der Reihe und wir sagen, dass $S_n \rightarrow S$ strebt. Ansonsten nennen wir die Reihe

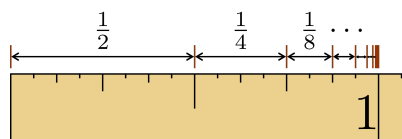


Abbildung 3.6: Veranschaulichung der (bei $k = 1$ beginnenden) geometrischen Reihe $\sum_{k=1}^{\infty} \frac{1}{2^k} = \frac{1}{2} + \frac{1}{4} + \frac{1}{8} + \dots = 1$. Die geometrische Folge $a_k = \frac{1}{2^k}$ konvergiert hinreichend schnell gegen 0, sodass die Reihe trotz unendlich vieler Summanden zum Grenzwert 1 konvergiert. Bildquelle: Wikipedia.

divergent.^a

^aDas Symbol $\sum_{k=1}^{\infty} a_k$ hat also zwei Bedeutungen: Zum einen bezeichnet es die unendliche Folge der Partialsummen, also die Reihe. Zum anderen – falls er existiert – deren Grenzwert.

Cauchy-Kriterium: Die Reihe $\sum_{k=1}^n a_k$ konvergiert genau dann, wenn sie eine Cauchy-Reihe ist, dh. wenn ihre Partialsummen eine Cauchy-Folge bilden. Dann gilt

$$\forall \varepsilon > 0 \exists N_{\varepsilon} \in \mathbb{N} \forall m > n \geq N_{\varepsilon} : |S_m - S_n| = |a_{n+1} + a_{n+2} + \dots + a_m| < \varepsilon. \quad (3.25)$$

Anders formuliert: Die Reihe konvergiert genau dann, wenn für alle $m > n$ und $n \rightarrow \infty$ gilt:

$$|S_m - S_n| = \left| \sum_{k=n+1}^m a_k \right| \rightarrow 0. \quad (3.26)$$

Insbesondere muss für Konvergenz der Fall $m = n + 1$ abgedeckt sein. Aus der Konvergenz folgt also $a_{n+1} \rightarrow 0$. Dieses sogenannte Nullfolgenkriterium beweisen wir später nochmal explizit.

Beispiel: Für $q \neq 1$ gilt^a

$$S_n := \sum_{k=0}^n q^k = \frac{1 - q^{n+1}}{1 - q}. \quad (3.27)$$

Beweis: Es ist $S_n = 1 + q + q^2 + \dots + q^n$. Daher ist $1 + qS_n = 1 + q + q^2 + \dots + q^{n+1} = S_n + q^{n+1}$. Daraus folgt $S_n = (1 - q^{n+1})/(1 - q)$. $\square$

^aDiese Folge startet mit dem Index $n = 0$, also beim Folgenglied S_0 .

Für $|q| < 1$ ist $\lim_{n \rightarrow \infty} q^{n+1} = 0$ und daher ist (für $m > n$) auch $\lim_{n \rightarrow \infty} |S_m - S_n| = 0$. Also konvergiert die *geometrische Reihe* $\sum_{k=0}^{\infty} q^k$ mit dem Wert

$$\sum_{k=0}^{\infty} q^k = \lim_{n \rightarrow \infty} S_n = \frac{1}{1 - q}. \quad (3.28)$$

Beispielsweise erhalten wir für $q = \frac{1}{2}$ den Wert $\sum_{k=0}^{\infty} \frac{1}{2^k} = 1 + \frac{1}{2} + \frac{1}{4} + \frac{1}{8} + \dots = 2$. Startet die Summe bei $k = 1$, so erhält man $\sum_{k=1}^{\infty} q^k = \sum_{k=0}^{\infty} q^k - 1 = \frac{1}{1 - q} - 1 = \frac{q}{1 - q}$. Figur 3.6 veranschaulicht diese Reihe für $q = \frac{1}{2}$.

Beispiel: Die harmonische Reihe

$$\sum_{k=1}^{\infty} \frac{1}{k} = 1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \dots \quad (3.29)$$

ist divergent. *Beweis:* Wir wählen ein beliebiges $n \in \mathbb{N}$ und betrachten die beiden Folgenglieder (Partialsummen) $S_n = 1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{n}$ sowie $S_{2n} = 1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{2n}$. Dann gilt für alle n :

$$S_{2n} - S_n = \frac{1}{n+1} + \frac{1}{n+2} + \dots + \frac{1}{2n} \geq \frac{n}{2n} = \frac{1}{2}. \quad (3.30)$$

Also ist $(S_n)_{n \in \mathbb{N}}$ keine Cauchy-Folge – denn dann müsste die Abweichung später Folgenglieder beliebig klein sein – und daher divergent. $\square$

Die harmonische Reihe strebt zu unbegrenzt großen Werten und hat den uneigentlichen Grenzwert ∞ . Wir sehen hier auf sehr eindrückliche Weise, dass es für die Konvergenz der (harmonischen) Reihe nicht genügt, dass die Glieder der (harmonischen) Folge $a_k = \frac{1}{k}$ gegen 0 streben. Formlos formuliert: Die Folgenglieder a_k streben zu langsam gegen 0, sodass die Partialsummen S_n (dh. die Folgenglieder der harmonischen Reihe) unbeschränkt bis ins Unendliche anwachsen (siehe Figur 3.7).

Konvergenzkriterien

Wie können wir feststellen, ob eine Reihe konvergiert oder nicht? Es gibt neben dem Cauchy-Kriterium eine Vielzahl an weiteren Kriterien. Wir besprechen im Folgenden einige wesentliche.

Nullfolgenkriterium

Definition: Eine Folge $(a_k)_{k \in \mathbb{N}}$ heißt *Nullfolge*, wenn sie gegen 0 konvergiert, also wenn gilt $\lim_{k \rightarrow \infty} a_k = 0$.

Satz (Nullfolgenkriterium): Gegeben eine Folge $(a_k)_{k \in \mathbb{N}}$ mit Partialsummen $S_n = \sum_{k=1}^n a_k$. Falls $(a_k)_{k \in \mathbb{N}}$ divergent ist oder $\lim_{n \rightarrow \infty} a_n \neq 0$, dann divergiert die Reihe $(S_n)_{n \in \mathbb{N}}$. Anders formuliert: Bildet die Folge der Summanden einer Reihe keine Nullfolge, dann divergiert die Reihe.

Beweis: Nehmen wir zunächst an, dass S_n gegen S konvergiert. Dann gilt, dass a_n gegen 0 konvergiert, also eine Nullfolge ist:

$$\lim_{n \rightarrow \infty} a_n = \lim_{n \rightarrow \infty} \left(\sum_{k=1}^n a_k - \sum_{k=1}^{n-1} a_k \right) = \lim_{n \rightarrow \infty} S_n - \lim_{n \rightarrow \infty} S_{n-1} = S - S = 0. \quad (3.31)$$

Dh. im Umkehrschluss^a, dass wenn a_n keine Nullfolge ist, S_n divergiert. $\square$

^a $(A \Rightarrow B) \Rightarrow (\neg B \Rightarrow \neg A)$

Das bedeutet, für Konvergenz von S_n ist $\lim_{n \rightarrow \infty} a_n = 0$ eine notwendige Voraussetzung. Es

ist keine hinreichende Voraussetzung. Die harmonische Reihe konvergiert nicht, obwohl die Folge $a_k = \frac{1}{k}$ eine Nullfolge ist. Symbolisch

$$\text{Konvergenz von } S_n \not\Rightarrow a_k \text{ ist Nullfolge} \quad (3.32)$$

Das Nullfolgenkriterium ist sehr intuitiv: Nur dann, wenn die Folgenglieder a_k gegen 0 streben, hat die Reihe S_n eine Chance zu konvergieren.

Beispiel: Die Reihe $\sum_{k=1}^{\infty} (-1)^k$ divergiert. *Beweis:* Der Grenzwert $\lim_{k \rightarrow \infty} (-1)^k$ existiert nicht. Damit ist $a_k = (-1)^k$ keine Nullfolge. Nach dem Nullfolgenkriterium muss die Reihe divergieren.

Beispiel: Die harmonische Reihe $\sum_{k=1}^{\infty} \frac{1}{k}$. Die Folge $a_k = \frac{1}{k}$ ist eine Nullfolge, da sie gegen 0 strebt. Das Nullfolgenkriterium kann daher keine Divergenz der Reihe begründen und lässt die Möglichkeit offen, dass die Reihe konvergiert (was aber in der Tat nicht der Fall ist).

Leibniz-Kriterium

Satz (Leibniz-Kriterium) (ohne Beweis): Sei $(a_k)_{k \in \mathbb{N}}$ eine monoton fallende Nullfolge. Dann konvergiert die alternierende Reihe

$$\sum_{k=0}^{\infty} (-1)^k a_k. \quad (3.33)$$

Beispiel: Die alternierende harmonische Reihe

$$\sum_{k=0}^{\infty} (-1)^k \frac{1}{k+1} = 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} \pm \dots \quad (3.34)$$

konvergiert, da $\frac{1}{k}$ eine monoton fallende Nullfolge ist. Der Grenzwert ist $\ln 2$.

Beispiel: Die Leibniz-Reihe

$$\sum_{k=0}^{\infty} (-1)^k \frac{1}{2k+1} = 1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} \pm \dots \quad (3.35)$$

konvergiert, da $\frac{1}{2k+1}$ eine monoton fallende Nullfolge ist. Der Grenzwert ist $\frac{\pi}{4}$.

Vergleichskriterium

Gegeben die beiden Folgen $(a_k)_{k \in \mathbb{N}}$ und $(b_k)_{k \in \mathbb{N}}$.

Satz (Majorantenkriterium): Wenn $\forall k \in \mathbb{N} : 0 \leq a_k \leq b_k$ und wenn $\sum_{k=1}^{\infty} b_k$ konvergiert,

dann konvergiert $\sum_{k=1}^{\infty} a_k$ mit Grenzwert

$$\sum_{k=1}^{\infty} a_k \leq \sum_{k=1}^{\infty} b_k. \quad (3.36)$$

Beweis: Wir definieren wie gewohnt die Partialsummen $S_n = \sum_{k=1}^n a_k$. Für alle $n \in \mathbb{N}$ gilt dann $S_n = \sum_{k=1}^n a_k \leq \sum_{k=1}^n b_k \leq \sum_{k=1}^{\infty} b_k$. Alle Partialsummen sind also nach oben beschränkt. Andererseits wachsen wegen $a_k \geq 0$ die Partialsummen. Die Folge der Partialsummen ist also monoton wachsend und beschränkt und muss daher wegen des Monotoniekriteriums konvergieren. Dh. die Reihe $\sum_{k=1}^{\infty} a_k$ muss konvergieren mit Grenzwert kleiner gleich $\sum_{k=1}^{\infty} b_k$. $\square$ Intuitiv: Die monoton wachsende Folge der Partialsummen $S_n = \sum_{k=1}^n a_k$ wird von oben durch den Grenzwert der Reihe $\sum_{k=1}^{\infty} b_k$ beschränkt und konvergiert daher.

Satz (Minorantenkriterium): Wenn $\forall k \in \mathbb{N} : 0 \leq a_k \leq b_k$ und wenn $\sum_{k=1}^{\infty} a_k$ divergiert, dann divergiert $\sum_{k=1}^{\infty} b_k$.

Beweis: Analog zum Majorantenkriterium. Intuitiv: Die unbeschränkt monoton wachsende Folge der Partialsummen $S_n = \sum_{k=1}^n a_k$ zwingt die Reihe $\sum_{k=1}^{\infty} b_k$, unbeschränkt zu sein.

Beispiel: Die Reihe

$$\sum_{k=1}^{\infty} \frac{1}{k^2} = 1 + \frac{1}{4} + \frac{1}{9} + \frac{1}{16} + \dots \quad (3.37)$$

konvergiert. *Beweis:* Zunächst zeigen wir:

$$\begin{aligned} S_n &:= \sum_{k=1}^n \frac{1}{k(k+1)} = \sum_{k=1}^n \left(\frac{1}{k} - \frac{1}{k+1} \right) = \left(\frac{1}{1} - \frac{1}{2} \right) + \left(\frac{1}{2} - \frac{1}{3} \right) + \dots + \left(\frac{1}{n} - \frac{1}{n+1} \right) \\ &= 1 - \frac{1}{n+1}. \end{aligned} \quad (3.38)$$

Die Summe in der ersten Zeile nennt man Teleskopsumme – eine endliche Summe von Differenzen, bei der je zwei Nachbarglieder (außer dem ersten und dem letzten) sich gegenseitig aufheben. Damit konvergiert $S_n \rightarrow 1$ für $n \rightarrow \infty$. Nun wählen wir $0 \leq a_k := \frac{1}{(k+1)^2} < b_k := \frac{1}{k(k+1)}$. Nach dem Majorantenkriterium muss die Reihe $\sum_{k=1}^{\infty} a_k$ konvergieren, weil $\sum_{k=1}^{\infty} b_k = 1$ konvergiert. Da $\sum_{k=1}^{\infty} \frac{1}{k^2} = 1 + \sum_{k=1}^{\infty} \frac{1}{(k+1)^2}$, konvergiert $\sum_{k=1}^{\infty} \frac{1}{k^2}$ auf einen Grenzwert im Intervall $[1, 2]$. $\square$ Ohne Beweis: Der Grenzwert von $\sum_{k=1}^{\infty} \frac{1}{k^2}$ ist $\frac{\pi^2}{6} \approx 1.6449$ (siehe Figur 3.7).

Beispiel: Die Reihe $\sum_{k=1}^{\infty} \frac{1}{\sqrt{k}}$ divergiert. *Beweis:* Es gilt $0 \leq a_k := \frac{1}{k} \leq b_k := \frac{1}{\sqrt{k}}$. Da die harmonische Reihe $\sum_{k=1}^{\infty} a_k$ divergiert, muss unsere Reihe $\sum_{k=1}^{\infty} b_k$ wegen des Minorantenkriteriums divergieren. $\square$

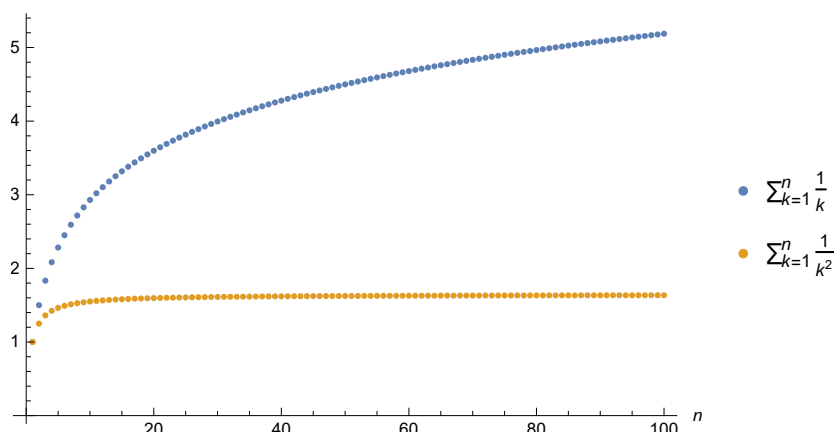


Abbildung 3.7: Die harmonische Reihe $\sum_{k=1}^{\infty} \frac{1}{k}$ (blau) ist divergent. Die Reihe $\sum_{k=1}^{\infty} \frac{1}{k^2}$ (orange) ist konvergent mit Limes $\frac{\pi^2}{6}$. Die Folge $a_k = \frac{1}{k}$ strebt nicht hinreichend schnell gegen 0, im Gegensatz zur Folge $a_k = \frac{1}{k^2}$.

Quotientenkriterium

Satz (Quotientenkriterium): Es sei $a_k \neq 0$ für alle $k \in \mathbb{N}$. (1) Falls

$$q := \lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| < 1, \quad (3.39)$$

dann konvergiert $\sum_{k=1}^{\infty} a_k$. (2) Falls $q > 1$, dann divergiert $\sum_{k=1}^{\infty} a_k$.

Beweis: (1) Wegen $0 \leq q < 1$ existiert ein $\bar{q} := \frac{q+1}{2}$, sodass $q < \bar{q} < 1$. Aufgrund von (3.39) gibt es ein $k_0 \in \mathbb{N}$, sodass $\forall k \geq k_0 : \left| \frac{a_{k+1}}{a_k} \right| < \bar{q}$. Daraus folgt

$$\left| \frac{a_k}{a_{k-1}} \right| \cdot \left| \frac{a_{k-1}}{a_{k-2}} \right| \cdot \left| \frac{a_{k-2}}{a_{k-3}} \right| \cdots \left| \frac{a_{k_0+1}}{a_{k_0}} \right| < \bar{q}^{k-k_0}. \quad (3.40)$$

Durch Kürzen und Umstellen erhalten wir

$$|a_k| < \frac{|a_{k_0}| \bar{q}^k}{\bar{q}^{n_0}} =: b_k. \quad (3.41)$$

Da die geometrische Reihe $\sum_{k=1}^{\infty} b_k$ konvergiert, muss nach dem Majorantenkriterium auch $\sum_{k=1}^{\infty} |a_k|$ konvergieren. Man spricht in diesem Fall von absoluter Konvergenz. Es gilt allgemein (ohne Beweis): Konvergiert $\sum_{k=1}^{\infty} |a_k|$, dann konvergiert $\sum_{k=1}^{\infty} a_k$. (2) Im Fall $q > 1$ kann a_k keine Nullfolge sein. Wegen des Nullfolgenkriteriums muss $\sum_{k=1}^{\infty} a_k$ divergieren. $\square$

Beispiel: Die Exponentialreihe

$$\exp(x) := \sum_{k=0}^{\infty} \frac{x^k}{k!} = 1 + x + \frac{x^2}{2!} + \frac{x^3}{3!} + \frac{x^4}{4!} + \dots \quad (3.42)$$

konvergiert für jedes $x \in \mathbb{R}$.^a *Beweis:* Für $x = 0$ ist $\exp(x) = 1$. Für $x \neq 0$ ist

$$\left| \frac{a_{k+1}}{a_k} \right| = \left| \frac{x^{k+1}}{(k+1)!} \cdot \frac{k!}{x^k} \right| = \frac{|x|}{k+1}. \quad (3.43)$$

Für $k \rightarrow \infty$ strebt $\left| \frac{a_{k+1}}{a_k} \right| \rightarrow 0 < 1$. Also konvergiert die Reihe wegen des Quotientenkriteriums. $\square$

^aEs konvergiert $\exp(z)$ sogar für jedes $z \in \mathbb{C}$.

Beispiel: Für die harmonische Reihe $\sum_{k=1}^{\infty} \frac{1}{k}$ erhalten wir $\left| \frac{a_{k+1}}{a_k} \right| = \frac{k}{k+1} \rightarrow 1$ für $k \rightarrow \infty$. Das Quotientenkriterium kann daher keine Aussage machen, ob diese Reihe konvergiert oder – was ja der Fall ist – divergiert.

Wurzelkriterium

Satz (Wurzelkriterium): Es sei $(a_k)_{k \in \mathbb{N}}$ eine Folge in $\mathbb{R}$ oder $\mathbb{C}$. Falls

$$q := \limsup_{k \rightarrow \infty} \sqrt[k]{|a_k|} < 1, \quad (3.44)$$

dann konvergiert $\sum_{k=1}^{\infty} a_k$ absolut. Falls $q > 1$, dann divergiert $\sum_{k=1}^{\infty} a_k$.^a

Beweis: Wenn $q < 1$, gilt mit $q < \bar{q} := \frac{q+1}{2} < 1$ für alle hinreichend späten Glieder $\sqrt[k]{|a_k|} < \bar{q} < 1$. Dh. $0 \leq |a_k| < \bar{q}^k < 1$. Da die geometrische Reihe $\sum_{k=0}^{\infty} \bar{q}^k$ konvergiert, muss wegen des Majorantenkriteriums auch die Reihe $\sum_{k=0}^{\infty} a_k$ absolut konvergieren. Für $q > 1$ setzen wir erneut $\bar{q} := \frac{1+q}{2}$, wobei nun $1 < \bar{q} < q$ gilt. Ab einem gewissen kritischen Folgenglied gilt für alle weiteren: $\sqrt[k]{|a_k|} > \bar{q} > 1$ und damit $1 < \bar{q}^k < |a_k|$. Da die geometrische Reihe $\sum_{k=0}^{\infty} \bar{q}^k$ divergiert, muss wegen des Minorantenkriteriums auch $\sum_{k=0}^{\infty} a_k$ divergieren. $\square$

^aFalls der Grenzwert existiert, kann man auch einfach den Limes anstelle des Limes Superior nehmen.

Das Wurzelkriterium ist schärfer als das Quotientenkriterium, kann also auch in Fällen eine Entscheidung liefern, in denen das Quotientenkriterium keine Aussage machen kann.

Beispiel: Wir betrachten die Folge $(a_k)_{k \in \mathbb{N}_0}$ mit $a_{2k} = \frac{1}{2^{2k}}$ und $a_{2k+1} = \frac{4}{2^{2k+1}}$, dh. $a = (1, 2, \frac{1}{4}, \frac{1}{2}, \frac{1}{16}, \frac{1}{8}, \frac{1}{64}, \dots)$. Hier ist $\left| \frac{a_{2k+1}}{a_{2k}} \right| = \frac{1}{8}$ und $\left| \frac{a_{2k+1+1}}{a_{2k+1}} \right| = 2$. Damit liefert das Quotientenkriterium keine Entscheidung. Aber wir können berechnen:

$$\lim_{k \rightarrow \infty} \sqrt[2k]{|a_{2k}|} = \lim_{k \rightarrow \infty} \sqrt[2k]{\frac{1}{2^{2k}}} = \frac{1}{2}, \quad (3.45)$$

$$\lim_{k \rightarrow \infty} \sqrt[2k+1]{|a_{2k+1}|} = \lim_{k \rightarrow \infty} \sqrt[2k+1]{\frac{4}{2^{2k+1}}} = \lim_{k \rightarrow \infty} \sqrt[2k+1]{4} \lim_{k \rightarrow \infty} \sqrt[2k+1]{\frac{1}{2^{2k+1}}} = \frac{1}{2}. \quad (3.46)$$

Damit konvergiert die Reihe $\sum_{k=0}^{\infty} a_k$ wegen des Wurzelkriteriums.

Beispiel: Wir betrachten die Reihe

$$\sum_{k=1}^{\infty} \frac{k^k}{c^k k!}. \quad (3.47)$$

Wir berechnen $\lim_{k \rightarrow \infty} \sqrt[k]{\left| \frac{k^k}{c^k k!} \right|} = \lim_{k \rightarrow \infty} \frac{k}{|c| \sqrt[k]{k!}} = \frac{e}{|c|}$. Hier haben wir verwendet, dass^a $\frac{k}{\sqrt[k]{k!}} \rightarrow e$ für $k \rightarrow \infty$. Nach dem Wurzelkriterium konvergiert die Reihe also für alle $|c| > e$ und divergiert für alle $|c| < e$.

^aDas ist äquivalent zur Stirlingschen Näherungsformel (für große k): $k! \approx \left(\frac{k}{e}\right)^k$

Beispiel: Wir betrachten die harmonische Reihe $\sum_{k=1}^{\infty} \frac{1}{k}$. Es gilt $\lim_{k \rightarrow \infty} \sqrt[k]{\frac{1}{k}} = \lim_{k \rightarrow \infty} \frac{1}{\sqrt[k]{k}}$. Betrachten wir nun $\sqrt[k]{k}$. Diese Folge hat den Grenzwert 1. Denn für jedes $\varepsilon > 0$ gibt es ein K_ε , sodass für alle $k \geq K_\varepsilon$ gilt: $|\sqrt[k]{k} - 1| < \varepsilon$. Die letzte Ungleichung lässt sich schreiben als $k < (1 + \varepsilon)^k$, und dies stimmt ab einem gewissen k immer, da die Funktion auf der rechten Seite exponentiell wächst, die linke Seite aber nur linear. Wegen $\lim_{k \rightarrow \infty} \sqrt[k]{\frac{1}{k}} = 1$ kann das Wurzelkriterium keine Aussage über die harmonische Reihe treffen.

Potenzreihen

Definition: Eine Potenzreihe in einer Variablen x ist eine (unendliche) Reihe der Form^a

$$f(x) = \sum_{k=0}^{\infty} a_k x^k. \quad (3.48)$$

^aStrenggenommen haben Potenzreihen die Form $f(x) = \sum_{n=0}^{\infty} a_n (x - x_0)^n$. Wir betrachten hier den Fall mit Entwicklungspunkt $x_0 = 0$.

Potenzreihen spielen eine grundlegende Rolle in der Analysis, insbesondere im Rahmen von Taylor-Reihen, die es ermöglichen, analytische Funktionen durch Potenzreihen darzustellen. In der Funktionentheorie erlauben Potenzreihen oft eine sinnvolle Fortsetzung reeller Funktionen in die komplexen Zahlen.

Zu beachten ist hier die Änderung der Notation im Vergleich zu den vorangegangenen Kapiteln. In der Potenzreihe $\sum_{k=0}^{\infty} a_k x^k$ heißen die summierten Folgenglieder nicht a_k sondern $a_k x^k$.

Konvergenzradius

Satz: Für jede Potenzreihe $f(x)$ gibt es einen *Konvergenzradius* $r \in \mathbb{R} \cup \{\infty\}$, sodass gilt:

- $f(x)$ ist konvergent für alle $|x| < r$,
- $f(x)$ ist divergent für alle $|x| > r$.

Beweis: Wegen des Wurzelkriteriums konvergiert die Potenzreihe $f(x)$, wenn gilt: $\limsup_{k \rightarrow \infty} \sqrt[k]{|a_k x^k|} = |x| \limsup_{k \rightarrow \infty} \sqrt[k]{|a_k|} < 1$. Mit $L := \limsup_{k \rightarrow \infty} \sqrt[k]{|a_k|}$ ist die Bedingung für Konvergenz also $|x| < \frac{1}{L} =: r$. Wenn $|x| > r$, dann divergiert die Reihe.



Wir unterscheiden drei Fälle:

$$r = \begin{cases} \frac{1}{L} & \text{falls } L \in \mathbb{R} \setminus 0, \\ \infty & \text{falls } L = 0, \\ 0 & \text{falls } L = \infty. \end{cases} \quad (3.49)$$

Falls $r = \frac{1}{L}$, dann konvergiert (divergiert) $f(x)$ für alle $|x| < \frac{1}{L}$ ($|x| > \frac{1}{L}$). Falls $r = \infty$, dann konvergiert $f(x)$ für alle x . Falls $r = 0$ ist, dann divergiert $f(x)$ für alle $x \neq 0$.

Oft ist es einfacher, das Quotientenkriterium heranzuziehen. Mit $L := \lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right|$ hat man wegen des Quotientenkriteriums dann Konvergenz, wenn $\left| \frac{a_{k+1}x^{k+1}}{a_kx^k} \right| = L|x| < 1$. Der Konvergenzradius ist $r := \frac{1}{L}$.

Beispiele: Wir listen ein paar bekannte Potenzreihen samt ihrem Konvergenzradius r auf:

$$\frac{1}{1-x} = \sum_{k=0}^{\infty} x^k = 1 + x + x^2 + x^3 + \dots \quad r = 1, \quad (3.50)$$

$$\exp(x) = \sum_{k=0}^{\infty} \frac{x^k}{k!} = 1 + x + \frac{x^2}{2!} + \frac{x^3}{3!} + \dots \quad r = \infty, \quad (3.51)$$

$$\sin(x) = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k+1}}{(2k+1)!} = x - \frac{x^3}{3!} + \frac{x^5}{5!} \mp \dots \quad r = \infty, \quad (3.52)$$

$$\cos(x) = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k}}{(2k)!} = 1 - \frac{x^2}{2!} + \frac{x^4}{4!} \mp \dots \quad r = \infty, \quad (3.53)$$

$$\ln(1+x) = \sum_{k=1}^{\infty} (-1)^{k+1} \frac{x^k}{k} = x - \frac{x^2}{2} + \frac{x^3}{3} \mp \dots \quad r = 1, \quad (3.54)$$

$$\sqrt{1+x} = 1 + \frac{1}{2}x - \frac{1}{2 \cdot 4}x^2 + \frac{1 \cdot 3}{2 \cdot 4 \cdot 6}x^3 \mp \dots \quad r = 1, \quad (3.55)$$

$$W(x) = \sum_{k=1}^{\infty} \frac{(-k)^{k-1} x^k}{k!} \quad r = \frac{1}{e}. \quad (3.56)$$

Die Lambertsche W -Funktion ist die Umkehrfunktion zu $x \exp(x)$, dh. $W(x \exp(x)) = x$.

Beispiel: Wir beweisen $r = 1$ für die Reihe $\frac{1}{1-x} = \sum_{k=0}^{\infty} x^k$ in (3.50). Dazu verwenden wir das Quotientenkriterium: $L = \lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| = \frac{1}{1} = 1$. Der Konvergenzradius ist daher $r = \frac{1}{L} = 1$. $\square$

Beispiel: Wir beweisen $r = \infty$ für die Reihe $\exp(x) = \sum_{k=0}^{\infty} \frac{x^k}{k!}$ in (3.51). Dazu verwenden wir das Quotientenkriterium: $L = \lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| = \lim_{k \rightarrow \infty} \frac{k!}{(k+1)!} = \lim_{k \rightarrow \infty} \frac{1}{k+1} = 0$. Der Konvergenzradius ist daher $r = \infty$. $\square$

Satz: Seien $f(x) = \sum_{k=0}^{\infty} f_k x^k$ und $g(x) = \sum_{k=0}^{\infty} g_k x^k$ zwei Potenzreihen mit Konvergenzradius r_f bzw. r_g . Setze $r := \min(r_f, r_g)$. Dann ist $h(x) = \sum_{k=0}^{\infty} (f_k + g_k) x^k$ eine Potenzreihe mit Konvergenzradius $\geq r$ und es gilt

$$h(x) = f(x) + g(x). \quad (3.57)$$

Beweis: Mittels Summensatz für Limites sowie Distributivgesetz:

$$\begin{aligned} f(x) + g(x) &= \lim_{n \rightarrow \infty} \sum_{k=0}^n f_k x^k + \lim_{n \rightarrow \infty} \sum_{k=0}^n g_k x^k \\ &= \lim_{n \rightarrow \infty} \left(\sum_{k=0}^n f_k x^k + \sum_{k=0}^n g_k x^k \right) \\ &= \lim_{n \rightarrow \infty} \sum_{k=0}^n (f_k + g_k) x^k = h(x). \quad \square \end{aligned} \quad (3.58)$$

Ohne Beweis: Für die Multiplikation gilt: $h(x) = \sum_{k=0}^{\infty} \sum_{l=0}^k f_l g_{k-l} x^k$ ist eine Potenzreihe mit Konvergenzradius $\geq r$ und es gilt $h(x) = f(x)g(x)$.

Die Exponentialreihe

Für die Exponentialreihe gilt das Additionstheorem (Beweis durch explizites Ausrechnen)

$$\exp(x+y) = \exp(x) \exp(y). \quad (3.59)$$

Mit $x = 1$ und $y = n - 1$ können wir schreiben: $\exp(n) = \exp(1) \exp(n-1)$. Und dann rekursiv weiter schreiben: $\exp(n) = \exp(1) \exp(n-1) = \exp(1) \exp(1) \exp(n-2)$ usw. Mithin:

$$\exp(n) = \exp(1)^n. \quad (3.60)$$

Wir setzen

$$e := \exp(1) = 2.718 \dots \quad (3.61)$$

wobei e Eulersche Zahl heißt.¹ Damit haben wir mithilfe von (3.60) somit $\exp(n) = e^n$ für $n \in \mathbb{N}$.

Dies können wir auf rationale Argumente erweitern: Die Zahl 1 können wir zerlegen in n Terme: $1 = \frac{1}{n} + \frac{1}{n} + \dots + \frac{1}{n}$. Somit gilt

$$e = \exp(1) = \exp\left(\frac{1}{n} + \frac{1}{n} + \dots + \frac{1}{n}\right) = \exp\left(\frac{1}{n}\right) \exp\left(\frac{1}{n}\right) \dots \exp\left(\frac{1}{n}\right) = \exp^n\left(\frac{1}{n}\right). \quad (3.62)$$

Wir ziehen die n -te Wurzel. Dies liefert $e^{\frac{1}{n}} = \exp\left(\frac{1}{n}\right)$. Nun nehmen wir die m -te Potenz,

¹Die Eulersche Zahl kann auch über den Grenzwert $e = \lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n$ definiert werden.

mit $m \in \mathbb{N}$, und bekommen:

$$e^{\frac{m}{n}} = \exp^m\left(\frac{1}{n}\right) = \underbrace{\exp \frac{1}{n} \exp \frac{1}{n} \cdots \exp \frac{1}{n}}_{m \text{ mal}} = \exp\left(\frac{m}{n}\right) \quad (3.63)$$

Da die rationalen Zahlen dicht in den reellen Zahlen liegen, gilt sogar für alle $x \in \mathbb{R}$:

$$e^x = \exp(x). \quad (3.64)$$

Die Exponentialfunktion ist sogar für alle $z \in \mathbb{C}$ konvergent. Für die imaginäre Zahl $z = i\varphi$ bekommen wir unter Verwendung der Potenzreihen für den Sinus (3.52) und Cosinus (3.53) die Darstellung für den Einheitskreis in der komplexen Ebene:

$$\begin{aligned} e^{i\varphi} &= \exp(i\varphi) = \sum_{k=0}^{\infty} \frac{(i\varphi)^k}{k!} \\ &= \sum_{k=0}^{\infty} \frac{(i\varphi)^{2k}}{(2k)!} + \sum_{k=0}^{\infty} \frac{(i\varphi)^{2k+1}}{(2k+1)!} \\ &= \sum_{k=0}^{\infty} (-1)^k \frac{\varphi^{2k}}{(2k)!} + i \sum_{k=0}^{\infty} (-1)^k \frac{\varphi^{2k+1}}{(2k+1)!} \\ &= \cos \varphi + i \sin \varphi. \end{aligned} \quad (3.65)$$

Hier wurde $i^{2k} = (i^2)^k = (-1)^k$ und $i^{2k+1} = i^1 \cdot i^{2k} = i(-1)^k$ verwendet. Damit haben wir über Potenzreihen die berühmte Eulersche Formel hergeleitet.

Mit der Eulerschen Formel können wir die Eulersche Identität formulieren:

$$e^{i\pi} = \cos \pi + i \sin \pi = -1. \quad (3.66)$$

Diese Identität kann man leicht umformen zu einer Gleichung, die bisweilen als die “schönste Formel der Mathematik” bezeichnet wird:

$$e^{i\pi} + 1 = 0 \quad (3.67)$$

Hier werden die fünf wichtigsten mathematischen Konstanten vereint: Das neutrale Element der Addition 0, das neutrale Element der Multiplikation 1, die Eulersche Zahl e , die Kreiszahl π sowie die imaginäre Einheit i .

4. Stetigkeit

Stetige Funktionen

Eine stetige Funktion (oder Abbildung) hat die Eigenschaft, dass hinreichend kleine Änderungen des Arguments x nur beliebig kleine Änderungen des Funktionswerts $f(x)$ verursachen. Vereinfacht formuliert: Stetige Funktionen kann man zeichnen, ohne den Stift bzw. die Kreide absetzen zu müssen.

Stetige Funktionen sind von fundamentaler Bedeutung in den Natur- und Ingenieurwissenschaften, weil viele relevante Größen stetige Funktionen von anderen Größen sind. Beispielsweise beschreibt man in der klassischen Physik die Bahn eines sich bewegenden Teilchens durch einen dreidimensionalen Ortsvektor, wobei alle Komponenten stetige Funktionen der Zeit t sind: $\mathbf{r}(t) = (x(t), y(t), z(t))$. Teilchen müssen sich kontinuierlich durch den Raum bewegen und können nicht von einem Ort an einen anderen springen. In der Quantenphysik werden Systeme durch eine (komplexe) Wellenfunktion $\psi(\mathbf{r}, t)$ beschrieben, die sich stetig in der Zeit entwickelt. Zu jedem Zeitpunkt t gibt der Absolutbetrag der Wellenfunktion die Wahrscheinlichkeit an, dass sich das Quantensystem am Ort $\mathbf{r}$ befindet.

Wir bleiben vorerst noch bei Funktionen von einer Variablen: Gegeben sei eine Funktion $f :]a, b[\rightarrow \mathbb{R}, x \mapsto f(x)$, wobei das offene Intervall $]a, b[$ auch ganz $\mathbb{R}$ sein kann. Wir interessieren uns für den Funktionswert an der Stelle x_0 .

Definition mittels Grenzwerts von Funktionen: Die Funktion f heißt *stetig* an der Stelle $x_0 \in]a, b[$, falls der Grenzwert $y_0 := \lim_{x \rightarrow x_0} f(x)$ existiert mit $y_0 = f(x_0)$.

Definition mittels Grenzwerts von Folgen: Die Funktion f hat an der Stelle x_0 den Grenzwert y_0 , falls für jede Nullfolge $(h_n)_{n \in \mathbb{N}}$ (dh. $\lim_{n \rightarrow \infty} h_n = 0$) gilt:

$$\lim_{n \rightarrow \infty} f(x_0 + h_n) = y_0. \quad (4.1)$$

Die Funktion f heißt *stetig* an der Stelle $x_0 \in]a, b[$, falls $y_0 = \lim_{n \rightarrow \infty} f(x_0 + h_n)$ existiert und der Grenzwert gleich dem Funktionswert an dieser Stelle ist: $y_0 = f(x_0) \in \mathbb{R}$.

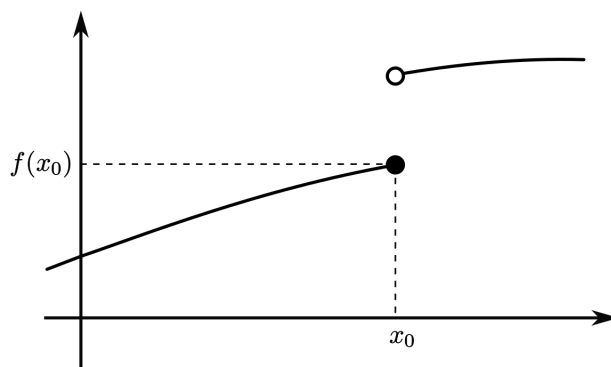


Abbildung 4.1: Diese Funktion $f(x)$ ist an der Stelle x_0 unstetig. Der Funktionswert an der Stelle x_0 ist $f(x_0)$, aber der Grenzwert $\lim_{x \rightarrow x_0} f(x)$ existiert nicht. Bildquelle: Wikibooks.

Die Funktion f heißt stetig im Intervall $]a, b[$, wenn sie stetig ist für alle $x_0 \in]a, b[$.

Formalisiert: f ist stetig im Intervall $]a, b[$, wenn

$$\forall x_0 \in]a, b[\forall (h_n)_{n \in \mathbb{N}} : \lim_{n \rightarrow \infty} h_n = 0 \Rightarrow \lim_{n \rightarrow \infty} f(x_0 + h_n) = f(x_0). \quad (4.2)$$

Figur 4.1 veranschaulicht eine an der Stelle x_0 unstetige Funktion.

Beispiel: Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x$ ist stetig in ganz $\mathbb{R}$. **Beweis:** Für alle $x_0 \in \mathbb{R}$ und für jede Nullfolge $(h_n)_{n \in \mathbb{N}}$ gilt

$$\lim_{n \rightarrow \infty} f(x_0 + h_n) = \lim_{n \rightarrow \infty} (x_0 + h_n) = x_0 + \lim_{n \rightarrow \infty} h_n = x_0 = f(x_0). \quad (4.3)$$

Beispiel: Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$ ist stetig in ganz $\mathbb{R}$. **Beweis:** Für alle $x_0 \in \mathbb{R}$ und für jede Nullfolge $(h_n)_{n \in \mathbb{N}}$ gilt

$$\begin{aligned} \lim_{n \rightarrow \infty} f(x_0 + h_n) &= \lim_{n \rightarrow \infty} (x_0 + h_n)^2 = \lim_{n \rightarrow \infty} (x_0^2 + 2x_0h_n + h_n^2) \\ &= x_0^2 + 2x_0 \lim_{n \rightarrow \infty} h_n + \left(\lim_{n \rightarrow \infty} h_n \right)^2 = x_0^2 = f(x_0). \end{aligned} \quad (4.4)$$

Beispiel: Die Funktion $f :]0, \infty[\rightarrow]0, \infty[$ mit $f(x) = \sqrt{x}$ ist stetig im ganzen Definitionsbereich $\mathbb{R}_{\geq 0}$. **Beweis:** In den Übungen.

Beispiel: Die Betragsfunktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = |x|$ ist stetig in ganz $\mathbb{R}$. **Beweis:** In den Übungen.

Beispiel: Die Vorzeichenfunktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = \operatorname{sgn}(x)$ ist stetig in $\mathbb{R} \setminus \{0\} =$

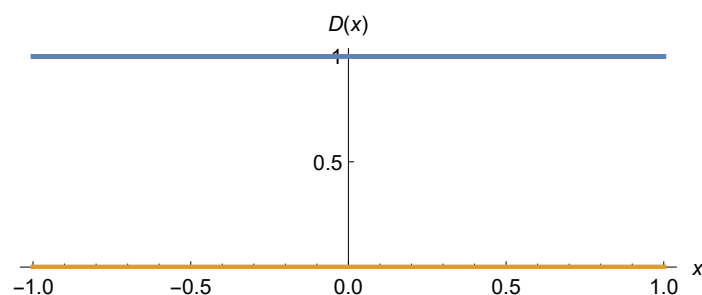


Abbildung 4.2: Die Dirichlet-Funktion $D(x)$ bildet alle rationalen Zahlen auf 1 ab (blau) und alle irrationalen Zahlen auf 0 (orange). Sie ist in jedem Punkt unstetig.

$] -\infty, 0[\cup]0, \infty[$. *Beweis:* Für alle $x_0 \neq 0$ gilt für alle Nullfolgen $(h_n)_{n \in \mathbb{N}}$:

$$\lim_{n \rightarrow \infty} f(x_0 + h_n) = \lim_{n \rightarrow \infty} \operatorname{sgn}(x_0 + h_n) = \operatorname{sgn}(x_0) = f(x_0). \quad (4.5)$$

Aber in $x_0 = 0$ führen zwei verschiedene Nullfolgen zu verschiedenen Limites:

$$h_n = -\frac{1}{n} : \lim_{n \rightarrow \infty} f(x_0 + h_n) = \lim_{n \rightarrow \infty} \operatorname{sgn}\left(-\frac{1}{n}\right) = -1, \quad (4.6)$$

$$h_n = +\frac{1}{n} : \lim_{n \rightarrow \infty} f(x_0 + h_n) = \lim_{n \rightarrow \infty} \operatorname{sgn}\left(+\frac{1}{n}\right) = +1. \quad (4.7)$$

Es gilt aber $f(x_0 = 0) = 0$. Die Nullfolge $h_n = 0$ würde das richtige Ergebnis bringen. Aber für Stetigkeit ist gefordert, dass *alle* Nullfolgen zum gleichen Limes $f(x_0)$ führen. Also ist $f(x)$ überall stetig außer an der Stelle 0. $\square$

Beispiel: Eine sogar in jedem Punkt unstetige Funktion ist die Dirichletsche Sprungfunktion (Figur 4.2)

$$D(x) = \begin{cases} 1 & \text{falls } x \text{ rational, dh. falls } x \in \mathbb{Q}, \\ 0 & \text{falls } x \text{ irrational, dh. falls } x \in \mathbb{R} \setminus \mathbb{Q}. \end{cases} \quad (4.8)$$

Dies ist ein sogenanntes *pathologisches Beispiel*, dh. eine Situation, die zwar mathematisch wohldefiniert ist, aber unserer Anschauung widerspricht und unerwünschte Eigenschaften aufweist.

Epsilon-Delta-Kriterium

Wir haben zwei Definitionen für die Stetigkeit einer Funktion f im Punkt x_0 kennengelernt, nämlich

1. Der Grenzwert $y_0 := \lim_{x \rightarrow x_0} f(x)$ existiert mit $y_0 = f(x_0)$.
2. Für jede Nullfolge $(h_n)_{n \in \mathbb{N}}$ existiert $y_0 := \lim_{n \rightarrow \infty} f(x_0 + h_n)$ mit $y_0 = f(x_0)$. Dies kann man auch so schreiben: Für jede Folge $(x_n)_{n \in \mathbb{N}}$ mit $\lim_{n \rightarrow \infty} x_n = x_0$ gilt $\lim_{n \rightarrow \infty} f(x_n) = f(x_0)$.

Wir fügen nun noch eine dritte Definition hinzu, nämlich das sogenannte

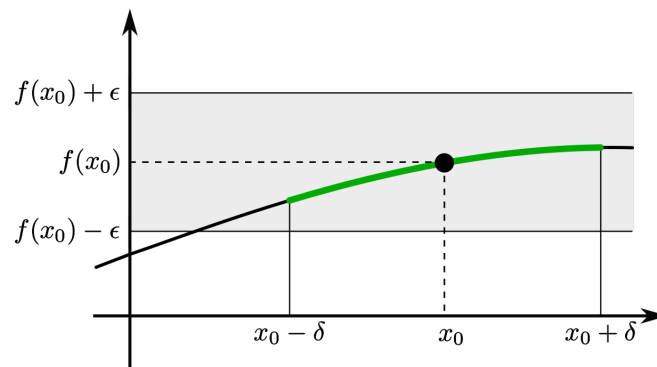


Abbildung 4.3: Veranschaulichung der ε - δ -Definition von Stetigkeit. Eine Funktion $f(x)$ ist genau dann stetig im Punkt x_0 , wenn für jede Fehlerschranke $\varepsilon > 0$ eine Schranke $\delta > 0$ gefunden werden kann, sodass alle Argumente in der δ -Umgebung von x_0 auf Funktionswerte in der ε -Umgebung von $f(x_0)$ abgebildet werden. Bildquelle: Wikibooks.

Epsilon-Delta-Kriterium: Die Funktion f ist stetig in x_0 , falls zu jeder Fehlerschranke $\varepsilon > 0$ ein $\delta > 0$ existiert, sodass für alle x mit $|x - x_0| < \delta$ gilt, dass $|f(x) - f(x_0)| < \varepsilon$.
Formal:

$$\forall \varepsilon > 0 \exists \delta > 0 : |x - x_0| < \delta \Rightarrow |f(x) - f(x_0)| < \varepsilon. \quad (4.9)$$

Intuitiv formuliert, bedeutet dies, dass zu jeder erlaubten Änderung des Funktionswertes eine hinreichend kleine Änderung im Argument gefunden werden kann, sodass man in der erlaubten Umgebung des Funktionswerts bleibt (Figur 4.3).

Beispiel: Die Funktion x ist stetig in ganz $\mathbb{R}$. **Beweis:** Wähle $\varepsilon > 0$. Dann soll es ein δ geben, sodass für alle x und x_0 mit $|x - x_0| < \delta$ gilt: $|f(x) - f(x_0)| = |x - x_0| = \delta < \varepsilon$. Mit jedem $\delta < \varepsilon$ ist das erfüllt, zB. also mit $\delta = \frac{\varepsilon}{2}$. $\square$

Beispiel: Die Funktion x^2 ist stetig in ganz $\mathbb{R}$. **Beweis:** Wähle $\varepsilon > 0$. Dann soll für alle x und x_0 mit $|x - x_0| < \delta$ gelten, dass $|f(x) - f(x_0)| < \varepsilon$. Wir berechnen:

$$\begin{aligned} |f(x) - f(x_0)| &= |x^2 - x_0^2| = |(x - x_0)(x + x_0)| = |x - x_0| |x + x_0| \\ &< \delta |x + x_0| = \delta |x - x_0 + 2x_0| \\ &\leq \delta (|x - x_0| + 2|x_0|) < \delta (\delta + 2|x_0|) = \delta^2 + 2|x_0|\delta. \end{aligned} \quad (4.10)$$

Von Zeile 2 auf 3 wurde die Dreiecksungleichung verwendet. Wir setzen $\delta^2 + 2|x_0|\delta = \varepsilon$ und lösen die quadratische Gleichung. Die beiden Lösungen sind $\delta_{\pm} = -|x_0| \pm \sqrt{|x_0|^2 + \varepsilon}$. Nur die Lösung $\delta = -|x_0| + \sqrt{|x_0|^2 + \varepsilon}$ ist positiv. Damit erhalten wir:

$$|f(x) - f(x_0)| < \delta^2 + 2|x_0|\delta = \varepsilon \quad (4.11)$$

Man beachte, dass δ nicht nur von ε abhängt, sondern auch von der Stelle x_0 . Aber man

kann für jedes ε an jeder Stelle x_0 das entsprechende $\delta = -|x_0| + \sqrt{|x_0|^2 + \varepsilon}$ angeben, sodass das Epsilon-Delta-Kriterium erfüllt ist. Also ist x^2 stetig in ganz $\mathbb{R}$. $\square$

Einseitige Grenzwerte

Nähert man sich der Stelle x_0 über Nullfolgen von der linken (bzw. rechten) Seite oder fordert man das Epsilon-Delta-Kriterium bloß für x -Werte, die links (bzw. rechts) von x_0 liegen, dann erhält man den linksseitigen (bzw. rechtsseitigen) Grenzwert

$$\lim_{x \rightarrow x_0^-} f(x), \quad (4.12)$$

$$\lim_{x \rightarrow x_0^+} f(x). \quad (4.13)$$

Wenn sowohl der links- als auch der rechtsseitige Grenzwert existieren und übereinstimmen, dann existiert der (beidseitige) Grenzwert, und umgekehrt.

Beispiel: Die Funktion $f(x) = x^2$ ist stetig in ganz $\mathbb{R}$ und hat an der Stelle x_0 den Grenzwert x_0^2 . Daher stimmen links- und rechtsseitiger Limes mit dem (beidseitigen) Limes überein: $\lim_{x \rightarrow x_0^-} f(x) = \lim_{x \rightarrow x_0^+} f(x) = \lim_{x \rightarrow x_0} f(x) = x_0^2$.

Beispiel: Die Funktion $f(x) = \operatorname{sgn}(x)$ hat an der Stelle $x_0 = 0$ den linksseitigen Grenzwert $\lim_{x \rightarrow 0^-} \operatorname{sgn}(x) = -1$ und den rechtsseitigen Grenzwert $\lim_{x \rightarrow 0^+} \operatorname{sgn}(x) = 1$. Spezielle Nullfolgen wurden in (4.6) und (4.7) gewählt, aber die Aussagen halten für alle Nullfolgen $h_n < 0$ bzw. $h_n > 0$. Da an der Stelle 0 links- und rechtsseitiger Limes nicht übereinstimmen, existiert an dieser Stelle kein Grenzwert und die Funktion ist an dieser Stelle unstetig.

Beispiel: Die Funktion $f(x) = \frac{1}{x}$ ist stetig in $\mathbb{R} \setminus \{0\}$. Für alle $x_0 \neq 0$ gilt für alle Nullfolgen $(h_n)_{n \in \mathbb{N}}$:

$$\lim_{n \rightarrow \infty} f(x_0 + h_n) = \lim_{n \rightarrow \infty} \frac{1}{x_0 + h_n} = \frac{1}{x_0} = f(x_0). \quad (4.14)$$

Also ist $f(x)$ stetig an allen Stellen $x_0 \neq 0$. An der Stelle $x_0 = 0$ hat $f(x)$ den linksseitigen und rechtsseitigen (uneigentlichen) Grenzwert

$$h_n < 0: \lim_{n \rightarrow \infty} f(x_0 + h_n) = \lim_{n \rightarrow \infty} \frac{1}{h_n} = -\infty, \quad (4.15)$$

$$h_n > 0: \lim_{n \rightarrow \infty} f(x_0 + h_n) = \lim_{n \rightarrow \infty} \frac{1}{h_n} = +\infty, \quad (4.16)$$

Der Limes an der Stelle $x_0 = 0$ existiert daher nicht und mithin ist $f(x)$ an dieser Stelle unstetig (Figur 4.4). Die Unstetigkeit an dieser Stelle folgt auch schon daraus, dass $f(0)$ undefiniert ist, also $x = 0$ nicht Teil der Definitionsmenge ist.

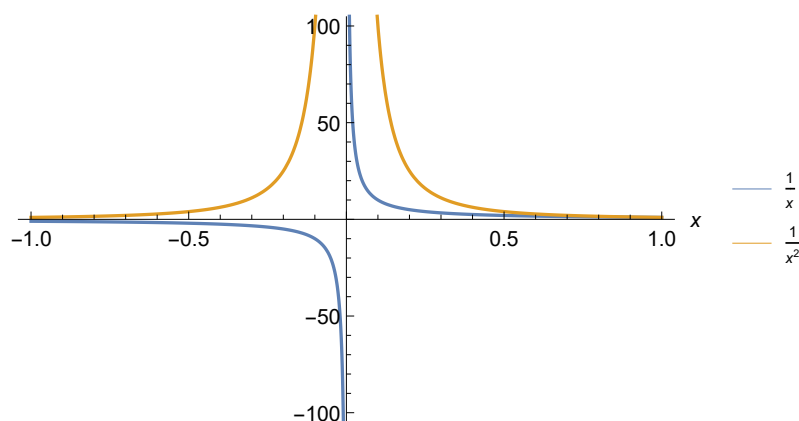


Abbildung 4.4: Die Funktion $f(x) = \frac{1}{x}$ (blau) ist stetig in $\mathbb{R} \setminus \{0\}$. Im Punkt $x_0 = 0$ ist der linksseitige Limes $-\infty$, während der rechtsseitige Limes $+\infty$ ist. An der Stelle 0 ist die Funktion daher unstetig und besitzt keinen Grenzwert, nicht einmal im uneigentlichen Sinn. Die Funktion $f(x) = \frac{1}{x^2}$ (orange) ist stetig in $\mathbb{R} \setminus \{0\}$. An der Stelle 0 ist sie zwar unstetig, besitzt aber den uneigentlichen Grenzwert $+\infty$.

Beispiel: Die Funktion $\frac{1}{x^2}$ ist stetig in $\mathbb{R} \setminus \{0\}$. Für alle $x_0 \neq 0$ und jede Nullfolge $(h_n)_{n \in \mathbb{N}}$ gilt

$$\lim_{n \rightarrow \infty} f(x_0 + h_n) = \lim_{n \rightarrow \infty} \frac{1}{(x_0 + h_n)^2} = \frac{1}{x_0^2} = f(x_0). \quad (4.17)$$

Also ist $f(x)$ stetig an allen Stellen $x_0 \neq 0$. An der Stelle $x_0 = 0$ stimmen der links- und rechtsseitige Limes überein:

$$h_n < 0: \lim_{n \rightarrow \infty} f(x_0 + h_n) = \lim_{n \rightarrow \infty} \frac{1}{h_n^2} = +\infty, \quad (4.18)$$

$$h_n > 0: \lim_{n \rightarrow \infty} f(x_0 + h_n) = \lim_{n \rightarrow \infty} \frac{1}{h_n^2} = +\infty, \quad (4.19)$$

An der Stelle $x_0 = 0$ existiert also der uneigentliche Grenzwert $+\infty$ (Figur 4.4). Dies ist aber keine reelle Zahl. Also ist die Funktion an der Stelle 0 unstetig. Die Unstetigkeit an dieser Stelle folgt auch schon daraus, dass $f(0)$ undefiniert ist, also $x = 0$ nicht Teil der Definitionsmenge ist.

Kombinationen stetiger Funktionen

Summen, Produkte, Quotienten

Satz: Summen, Produkte und Quotienten von stetigen Funktionen sind wieder stetig. Seien $f(x)$ und $g(x)$ stetige Funktionen in $]a, b[$. Dann gilt

1. $f + g$ ist stetig in $]a, b[$.
2. $f \cdot g$ ist stetig in $]a, b[$.
3. Wenn $\forall x \in]a, b[: g(x) \neq 0$, dann ist $\frac{f}{g}$ stetig in $]a, b[$.

Beweis von 1. Wir benennen $F := f + g$. Für alle Nullfolgen $(h_n)_{n \in \mathbb{N}}$ gilt (mithilfe des Summensatzes für Grenzwerte):

$$\begin{aligned}\lim_{n \rightarrow \infty} F(x_0 + h_n) &= \lim_{n \rightarrow \infty} [f(x_0 + h_n) + g(x_0 + h_n)] \\ &= \lim_{n \rightarrow \infty} f(x_0 + h_n) + \lim_{n \rightarrow \infty} g(x_0 + h_n) \\ &= f(x_0) + g(x_0) = F(x_0).\end{aligned}\tag{4.20}$$

Also ist $f + g$ stetig für alle x_0 in $]a, b[$ mit Grenzwert $f(x_0) + g(x_0)$. $\square$ Die Punkte 2. und 3. können auf ähnliche Weise bewiesen werden.

Mit diesem Satz folgt aus der Stetigkeit der Funktion $g(x) = x$ sowie der Stetigkeit aller konstanten Funktionen $c_k(x) = a_k$ die Stetigkeit aller Potenzreihen

$$\begin{aligned}f(x) &= c_0(x) + c_1(x)g(x) + c_2(x)g(x)g(x) + c_3(x)g(x)g(x)g(x) + \dots \\ &= a_0 + a_1x + a_2x^2 + a_3x^3 + \dots = \sum_{k=0}^{\infty} a_k x^k\end{aligned}\tag{4.21}$$

Verkettungen und Umkehrfunktionen

Satz (ohne Beweis): Verkettungen von stetigen Funktionen sind wieder stetig. Sei $f(x) : X \rightarrow Y$ stetig, und sei $h : Y \rightarrow Z$ stetig. Dann ist $h \circ f : X \rightarrow Z$ stetig.

Beispiel: Die Funktionen $\sin(x)$ und x^3 sind stetig in ganz $\mathbb{R}$. Also sind auch die Funktionen $\sin^3(x)$ und $\sin(x^3)$ stetig in ganz $\mathbb{R}$.

Satz (ohne Beweis): Umkehrfunktionen von stetigen (bijektiven) Funktionen sind wieder stetig.

Beispiel: Die Exponentialfunktion $\exp(x)$ ist stetig in ganz $\mathbb{R}$. Also ist auch die Umkehrfunktion $\ln(x)$ stetig in ganz $\mathbb{R}$.

Der Zwischenwertsatz

Der Zwischenwertsatz besagt, dass eine reelle Funktion f , die auf einem abgeschlossenen Intervall $[a, b]$ stetig ist, jeden Wert zwischen $f(a)$ und $f(b)$ mindestens einmal annimmt (Figur 4.5). Formal:

Zwischenwertsatz: Sei $-\infty < a < b < \infty$ und $f : [a, b] \rightarrow \mathbb{R}$ stetig. Dann gibt es zu jedem $s \in [\min(f(a), f(b)), \max(f(a), f(b))]$ ein $x \in [a, b]$ mit $f(x) = s$.

Beweis: Der Beweis des Zwischenwertsatzes erfolgt konstruktiv über das sogenannte Bisektionsverfahren. Wir nehmen an, dass $f(a) \leq f(b)$ ist, also gilt $f(a) \leq s \leq f(b)$;

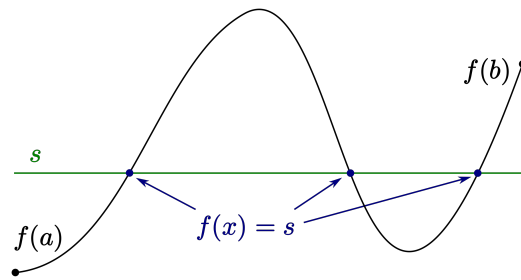


Abbildung 4.5: Zwischenwertsatz: Sei $f(x)$ eine auf $[a, b]$ stetige Funktion. Dann wird jeder Wert s , für den gilt $\min(f(a), f(b)) \leq s \leq \max(f(a), f(b))$, von der Funktion mindestens einmal erreicht, dh. es gibt mindestens ein $x \in [a, b]$, sodass $f(x) = s$. Bildquelle: Wikipedia.

der andere Fall $f(b) \leq s \leq f(a)$ funktioniert ganz analog. Wir setzen zu Beginn

$$a_1 := a, \quad b_1 := b, \quad x_1 := \frac{a_1 + b_1}{2}, \quad (4.22)$$

und berechnen $y_1 := f(x_1)$, also den Funktionswert in der Mitte.

- Falls $y_1 = s$, ist x_1 eine Lösung.
- Falls $y_1 < s$, setze $a_2 := x_1$ und $b_2 := b_1$.
- Falls $y_1 > s$, setze $a_2 := a_1$ und $b_2 := x_1$.

Nach diesem ersten Schritt hat man entweder eine Lösung gefunden (erster Fall) oder den Intervall halbiert (zweiter oder dritter Fall). Im zweiten Schritt halbiert man wieder das Intervall, prüft und versetzt eine Grenze. Und so weiter und so fort.

Allgemein: $x_n := \frac{a_n + b_n}{2}$. Falls $x_n = s$, ist eine Lösung gefunden. Falls $f(x_n) < s$, setze $a_{n+1} := x_n$ und $b_{n+1} := b_n$. Falls $f(x_n) > s$, setze $a_{n+1} := a_n$ und $b_{n+1} := x_n$. Die Folgen (a_n) bzw. (b_n) sind monoton wachsend bzw. fallend und beschränkt. Daher sind sie nach dem Monotoniekriterium konvergent, dh. es existiert

$$\lim_{n \rightarrow \infty} a_n = \alpha, \quad (4.23)$$

$$\lim_{n \rightarrow \infty} b_n = \beta. \quad (4.24)$$

Andererseits wissen wir, dass die Distanz der Folgenglieder exponentiell fällt:

$$b_n - a_n = \frac{b - a}{2^{n-1}}. \quad (4.25)$$

Das impliziert:

$$\lim_{n \rightarrow \infty} (b_n - a_n) = \lim_{n \rightarrow \infty} (b_n) - \lim_{n \rightarrow \infty} a_n = \beta - \alpha = 0. \quad (4.26)$$

Mithin ist $\alpha = \beta$. Nun müssen wir aber noch zeigen, dass an dieser Stelle $x = \alpha = \beta$ der Funktionswert s erreicht (und nicht etwa unstetig übersprungen) wird. Der links- und

rechtsseitige Limes von f sind:

$$\lim_{n \rightarrow \infty} f(a_n) \leq s, \quad (4.27)$$

$$\lim_{n \rightarrow \infty} f(b_n) \geq s. \quad (4.28)$$

Wegen der Stetigkeit von f müssen diese beiden Grenzwerte übereinstimmen und zum Funktionswert an dieser Stelle $x = \alpha = \beta$ führen. Es muss also gelten:

$$f(\alpha) = f(\beta) = \lim_{n \rightarrow \infty} f(a_n) = \lim_{n \rightarrow \infty} f(b_n) = s. \quad (4.29)$$

Mithin wird der Wert s erreicht. $\square$

Der letzte Schritt im Beweis ist essentiell. Nehmen wir zur Veranschaulichung die Signumfunktion $f(x) = \operatorname{sgn}(x)$ im Intervall $[-1, 1]$ und $s = \frac{1}{2}$. Das Bisektionsverfahren liefert $\alpha = \beta = 0$. Aber der links- bzw. rechtsseitige Limes sind $\lim_{n \rightarrow \infty} f(a_n) = -1$ bzw. $\lim_{n \rightarrow \infty} f(b_n) = +1$. Dh. die Funktion ist an dieser Stelle unstetig. Der Wert s wird nicht erreicht. Die Unstetigkeit der Signumfunktion hebt den Zwischenwertsatz aus, der nur für stetige Funktionen gilt.

Eine anschauliche Animation des Bisektionsverfahrens (mit leicht anderer Notation) findet man *hier*.

5. Differentialrechnung

Die Aufgabe der Differentialrechnung ist die Untersuchung lokaler Veränderungen von Funktionen. Während eine stetige Funktion ihren Argumenten kontinuierlich Bildwerte zuordnet, zeigt die Differentialrechnung, wie stark dieser Zusammenhang ist, also wie stark sich der Bildwert ändert, wenn man das Argument variiert. Anschaulich ist die Aufgabenstellung der Differentialrechnung also die Bestimmung der Steigung oder der Tangente in einem Punkt einer beliebigen Funktion.

Dieses sogenannte *Tangentenproblem* reicht zurück bis in den Anfang des 17. Jahrhunderts. Erst Ende des 17. Jahrhunderts gelang der Durchbruch: Isaac Newton und Gottfried Wilhelm Leibniz (Figur 1) lösten das Problem unabhängig voneinander und mit unterschiedlichen Ansätzen. Newton verwendete einen physikalischen Ansatz über die Momentangeschwindigkeit von Körpern, Leibniz einen geometrischen.

5.1 In einer Variablen

Differenzenquotient

Wir betrachten eine Funktion f und interessieren uns für die *Steigung* an der Stelle x_0 . Dazu betrachten wir die Funktion an der Stelle x_0 und an der Stelle $x_0 + \Delta x$, mit einer "Schrittweite" $\Delta x > 0$. Durch die beiden Punkte $(x_0, f(x_0))$ und $(x_0 + \Delta x, f(x_0 + \Delta x))$ können wir eine Gerade legen, genannt Sekante. Der Differenzenquotient ist die Steigung dieser Sekante (Figur 5.1):

$$S_{\Delta} = \frac{\Delta y}{\Delta x} = \frac{f(x_0 + \Delta x) - f(x_0)}{(x_0 + \Delta x) - x_0} = \frac{f(x_0 + \Delta x) - f(x_0)}{\Delta x}. \quad (5.1)$$

Strenggenommen ist das der sogenannte Vorwärts-Differenzenquotient. Der Rückwärts-Differenzenquotient verwendet die Sekante durch die Punkte $(x_0 - \Delta x, f(x_0 - \Delta x))$ und $(x_0, f(x_0))$: $\frac{f(x_0) - f(x_0 - \Delta x)}{\Delta x}$. Der zentrale Differenzenquotient verwendet die beiden Punkte $(x_0 - \Delta x, f(x_0 - \Delta x))$ und $(x_0 + \Delta x, f(x_0 + \Delta x))$: $\frac{f(x_0 + \Delta x) - f(x_0 - \Delta x)}{2\Delta x}$.

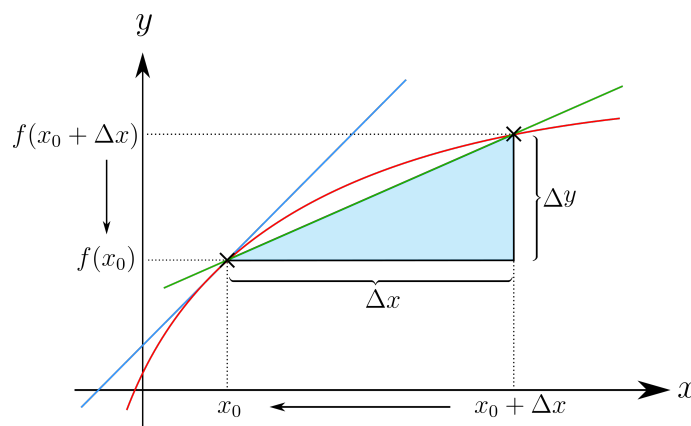


Abbildung 5.1: Der Vorwärts-Differenzenquotient einer Funktion f (rot) an der Stelle x_0 ist die Steigung der Sekante (grün), die durch die Punkte $(x_0, f(x_0))$ und $(x_0 + \Delta x, f(x_0 + \Delta x))$ geht, wobei $\Delta x > 0$. Wenn Δx gegen 0 geht, nähert sich die Sekante der Tangente (blau) an. Bildquelle: Wikipedia.

Differenzierbarkeit

Die Steigung der Funktion in einem Punkt x_0 ist durch die Steigung der *Tangente* in diesem Punkt gegeben. Die Tangente erhält man, wenn man für die Sekante Δx gegen 0 gehen lässt. Dies führt zu folgender

Definition (Differenzierbarkeit): Sei $f :]a, b[\rightarrow \mathbb{R}$ mit $a, b \in \mathbb{R} \cup \{-\infty, +\infty\}$. Die Funktion f heißt *differenzierbar* an der Stelle $x_0 \in]a, b[$, wenn für alle Nullfolgen $(h_n)_{n \in \mathbb{N}}$ mit $h_n \neq 0$ der eigentliche Grenzwert

$$\lim_{n \rightarrow \infty} \frac{f(x_0 + h_n) - f(x_0)}{h_n} \quad (5.2)$$

existiert, dh. eine reelle Zahl ist. Dann schreibt man für diesen Grenzwert $f'(x_0)$ und nennt ihn die erste Ableitung von f in x_0 .

Man nennt f differenzierbar auf dem Intervall $]a, b[$, wenn f in jedem Punkt $x_0 \in]a, b[$ differenzierbar ist. Die Funktion $f' :]a, b[\rightarrow \mathbb{R}$, $x \mapsto f'(x)$ ist die Ableitung (oder Ableitungsfunktion) von f .

Beispiel: Die Funktion $f(x) = x^2$ ist differenzierbar auf ganz $\mathbb{R}$. *Beweis:* Für jedes $x_0 \in \mathbb{R}$ und jede Nullfolge $(h_n)_{n \in \mathbb{N}}$ mit $h_n \neq 0$ gilt

$$\begin{aligned} \lim_{n \rightarrow \infty} \frac{f(x_0 + h_n) - f(x_0)}{h_n} &= \lim_{n \rightarrow \infty} \frac{x_0^2 + 2x_0h_n + h_n^2 - x_0^2}{h_n} \\ &= \lim_{n \rightarrow \infty} (2x_0 + h_n) = 2x_0 =: f'(x_0). \quad \square \end{aligned} \quad (5.3)$$

Ohne Beweis: Allgemein gilt, dass $f(x) = x^p$ mit $p \in \mathbb{R}$ auf ganz $\mathbb{R}$ differenzierbar ist mit der Ableitungsfunktion $f'(x) = px^{p-1}$. Figur 5.2 zeigt die Graphen für die Funktion $f(x) = x^3$ und ihre Ableitung $f'(x) = 3x^2$.

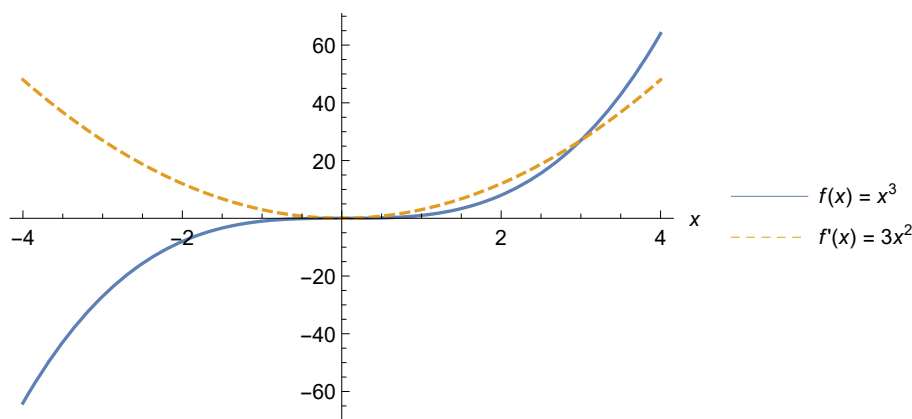


Abbildung 5.2: Die Funktion $f(x) = x^3$ ist differenzierbar auf ganz $\mathbb{R}$. Die Ableitungsfunktion $f'(x) = 3x^2$ hat an jeder Stelle x_0 genau den Wert, der der Steigung der Tangente der ursprünglichen Funktion $f(x)$ an der Stelle x_0 entspricht. Beispielsweise hat an der Stelle $x_0 = 4$ die Steigung der Tangente im Punkt $(4, f(4) = 64)$ den Wert $f'(4) = 48$.

Satz (ohne Beweis): Aus der Differenzierbarkeit einer Funktion (in einem Punkt bzw. Intervall) folgt deren Stetigkeit (an dieser Stelle bzw. auf diesem Intervall).

Im Umkehrschluss ist daher eine nicht stetige Funktion auch nicht differenzierbar. Andererseits muss eine stetige Funktion aber nicht notwendigerweise differenzierbar sein:

$$\text{Differenzierbarkeit} \xrightarrow{\neq} \text{Stetigkeit} \quad (5.4)$$

Differenzierbarkeit ist also das “stärkere” Kriterium. Das sehen wir anhand von folgendem

Beispiel: Die Betragsfunktion $f(x) = |x|$ ist, wie wir schon wissen, stetig in ganz $\mathbb{R}$. Aber sie ist nicht differenzierbar im Punkt $x_0 = 0$. Verschiedene Nullfolgen führen zu verschiedenen Grenzwerten:

$$h_n = -\frac{1}{n} : \lim_{n \rightarrow \infty} \frac{|x_0 + h_n| - |x_0|}{h_n} = \lim_{n \rightarrow \infty} \frac{+1/n}{-1/n} = -1, \quad (5.5)$$

$$h_n = +\frac{1}{n} : \lim_{n \rightarrow \infty} \frac{|x_0 + h_n| - |x_0|}{h_n} = \lim_{n \rightarrow \infty} \frac{+1/n}{+1/n} = +1. \quad (5.6)$$

Wenn wir uns $x_0 = 0$ von links bzw. rechts annähern, ist die Steigung -1 bzw. $+1$. An der Stelle $x_0 = 0$ gibt es daher keinen Grenzwert. Die Betragsfunktion ist differenzierbar nur auf $\mathbb{R} \setminus \{0\}$. Dort, also für alle $x \neq 0$, ist die Ableitungsfunktion $f'(x) = |x|' = \frac{x}{|x|} = \text{sgn}(x)$.

Kombinationen differenzierbarer Funktionen

Summen, Produkte, Quotienten

Satz: Seien $f(x)$ und $g(x)$ differenzierbare Funktionen auf $]a, b[$. Dann gelten:

1. Summenregel: $f + g$ ist differenzierbar auf $]a, b[$ mit $(f + g)'(x) = f'(x) + g'(x)$.
2. Produktregel: $f \cdot g$ ist differenzierbar auf $]a, b[$ mit $(fg)'(x) = f'(x)g(x) + f(x)g'(x)$.
3. Quotientenregel: Wenn $\forall x \in]a, b[: g(x) \neq 0$, dann ist $\frac{f}{g}$ differenzierbar auf $]a, b[$ mit $\left(\frac{f}{g}\right)'(x) = \frac{f'(x)g(x) - f(x)g'(x)}{g^2(x)}$.

Beweis von 1. Mittels Sommensatz von Grenzwerten:

$$\begin{aligned}
 (f + g)'(x_0) &= \lim_{n \rightarrow \infty} \frac{(f + g)(x_0 + h_n) - (f + g)(x_0)}{h_n} \\
 &= \lim_{n \rightarrow \infty} \frac{f(x_0 + h_n) + g(x_0 + h_n) - [f(x_0) + g(x_0)]}{h_n} \\
 &= \lim_{n \rightarrow \infty} \frac{f(x_0 + h_n) - f(x_0)}{h_n} + \lim_{n \rightarrow \infty} \frac{g(x_0 + h_n) - g(x_0)}{h_n} \\
 &= f'(x_0) + g'(x_0).
 \end{aligned} \tag{5.7}$$

Beweis von 2. Mittels Produktsatz, Sommensatz und Stetigkeit:

$$\begin{aligned}
 (fg)'(x_0) &= \lim_{n \rightarrow \infty} \frac{(fg)(x_0 + h_n) - (fg)(x_0)}{h_n} \\
 &= \lim_{n \rightarrow \infty} \frac{f(x_0 + h_n)g(x_0 + h_n) - f(x_0)g(x_0)}{h_n} \\
 &= \lim_{n \rightarrow \infty} \left(\frac{f(x_0 + h_n)g(x_0 + h_n) - f(x_0)g(x_0 + h_n)}{h_n} + \frac{f(x_0)g(x_0 + h_n) - f(x_0)g(x_0)}{h_n} \right) \\
 &= \lim_{n \rightarrow \infty} \frac{f(x_0 + h_n) - f(x_0)}{h_n} \lim_{n \rightarrow \infty} g(x_0 + h_n) + \lim_{n \rightarrow \infty} f(x_0) \frac{g(x_0 + h_n) - g(x_0)}{h_n} \\
 &= f'(x_0)g(x_0) + f(x_0)g'(x_0).
 \end{aligned} \tag{5.8}$$

Hier wurde die Stetigkeit von g verwendet, dh. $\lim_{n \rightarrow \infty} g(x_0 + h_n) = g(x_0)$.

Mit diesem Satz können wir nun ganz ähnlich wie im Kapitel zur Stetigkeit für Kombinationen von differenzierbaren Funktionen zeigen, dass sie differenzierbar sind.

Beispiel: Wir betrachten die auf ganz $\mathbb{R}$ differenzierbare Polynomfunktion $f(x) = 3x^2 + 2x + 1$. Mit der Summen- und Produktregel erhalten wir die erste Ableitung:

$$\begin{aligned}
 f'(x) &= (3x^2 + 2x + 1)' \\
 &= (3x^2)' + (2x)' + (1)' \\
 &= (3)'x^2 + 3(x^2)' + (2)'x + 2(x)' + (1)' \\
 &= 6x + 2.
 \end{aligned} \tag{5.9}$$

Beispiel: Gegeben eine Potenzreihe $f(x) = \sum_{k=0}^{\infty} a_k x^k$ mit Konvergenzradius r . Die Funktion $f(x)$ ist differenzierbar auf $] -r, r[$. Die erste Ableitung ist

$$f'(x) = \sum_{k=1}^{\infty} a_k k x^{k-1} \tag{5.10}$$

und hat wieder den Konvergenzradius r .

Beispiel: Die Exponentialfunktion $e^x = \exp(x) = \sum_{k=0}^{\infty} \frac{x^k}{k!}$ ist differenzierbar auf ganz $\mathbb{R}$. Die erste Ableitung ist

$$\exp'(x) = \left(\sum_{k=0}^{\infty} \frac{x^k}{k!} \right)' = \sum_{k=1}^{\infty} \frac{kx^{k-1}}{k!} = \sum_{k=1}^{\infty} \frac{x^{k-1}}{(k-1)!} = \sum_{k=0}^{\infty} \frac{x^k}{k!} = \exp(x), \quad (5.11)$$

also die Funktion selbst.

Beispiel: Die Ableitung der Sinus-Funktion (3.52) ist

$$\begin{aligned} \sin'(x) &= \left(\sum_{k=0}^{\infty} (-1)^k \frac{x^{2k+1}}{(2k+1)!} \right)' = \sum_{k=0}^{\infty} (-1)^k \frac{(2k+1)x^{2k}}{(2k+1)!} = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k}}{(2k)!} \\ &= \cos(x). \end{aligned} \quad (5.12)$$

Beispiel: Ableitung der Cosinus-Funktion via Potenzreihe. *Lösung:* In den Übungen.

Verkettungen und Umkehrfunktionen

Satz (Kettenregel) (ohne Beweis): Die Funktionen $f :]a, b[\rightarrow]c, d[$ und $g :]c, d[\rightarrow \mathbb{R}$ seien differenzierbar auf ihren jeweiligen Definitionsbereichen. Dann ist $g \circ f$ differenzierbar auf $]a, b[$ und es gilt dort

$$(g \circ f)'(x) = [g(f(x))]' = g'(f(x)) f'(x). \quad (5.13)$$

Man bezeichnet $g'(f(x))$ als die äußere Ableitung und $f'(x)$ als die innere Ableitung.

Beispiel: Wir betrachten die Funktionen $f(x) = x^2$ und $g(x) = \sin x$. Die Funktion $g(f(x)) = \sin x^2$ hat die Ableitung

$$[g(f(x))]' = (\sin x^2)' = (\cos x^2) \cdot (2x) = 2x \cos x^2. \quad (5.14)$$

Wenn mehr als zwei Funktionen verkettet sind, wird die Kettenregel rekursiv angewandt. Sei $f(x) = u(v(w(x)))$. Dann gilt:

$$f'(x) = u'(v(w(x))) \cdot (v(w(x)))' = u'(v(w(x))) \cdot v'(w(x)) \cdot w'(x). \quad (5.15)$$

Beispiel: Wir betrachten die Funktion $f(x) = \sqrt{\cos x^3}$. Die Ableitungsfunktion ist gegeben durch

$$f'(x) = \frac{1}{2} (\cos x^3)^{-\frac{1}{2}} \cdot (-\sin x^3) \cdot 3x^2 = -\frac{3x^2 \sin x^3}{2\sqrt{\cos x^3}}. \quad (5.16)$$

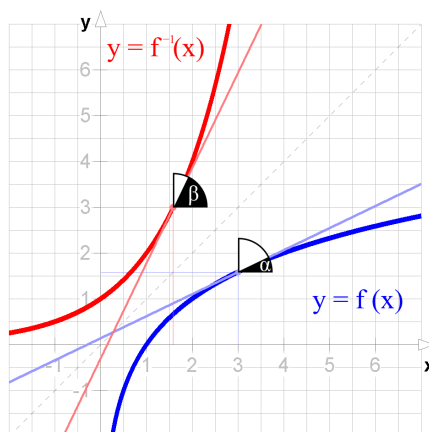


Abbildung 5.3: Umkehrregel: Gezeichnet sind die (bijektive) Funktion $f(x)$ in blau und deren Umkehrfunktion $f^{-1}(x)$ in rot, die man durch Spiegelung von f an der Diagonalen (grau strichliert) erhält. Wir interessieren uns für die Ableitung der Inversen f^{-1} an der Stelle $x_0 \approx 1.55$, also für $(f^{-1})'(x_0) = \tan \beta$, dh. für die Steigung der Tangente (dünne rote Linie). An dieser Stelle x_0 hat die Inverse den Wert $f^{-1}(x_0) = 3$. Nun betrachten wir die ursprüngliche Funktion an dieser Stelle $f^{-1}(x_0) = 3$, also $f(f^{-1}(x_0)) = f(3)$. Die Tangente an dieser Stelle (dünne blaue Linie) hat die Steigung $f'(f^{-1}(x_0)) = \tan \alpha$. Aufgrund der Spiegelsymmetrie gilt $\alpha + \beta = \frac{\pi}{2}$. Daher ist $\tan \beta = \tan(\frac{\pi}{2} - \alpha) = \frac{1}{\tan \alpha}$. Das ist genau die Umkehrregel an der Stelle x_0 : $(f^{-1})'(x_0) = \frac{1}{f'(f^{-1}(x_0))}$. Dies gilt für alle Stellen x_0 . Bildquelle: Wikipedia.

Satz (Umkehrregel): Wenn f differenzierbar ist, ist auch die inverse Funktion f^{-1} differenzierbar, sofern f' keine Nullstelle im Ableitungsbereich hat. Die Ableitung der inversen Funktion ist dann gegeben durch

$$(f^{-1})'(x) = \frac{1}{f'(f^{-1}(x))}. \quad (5.17)$$

Beweisidee: Figur 5.3 zeigt eine graphische Veranschaulichung der Umkehrregel anhand eines konkreten Beispiels

Beispiel: Gegeben $f(x) = x^2$. Dann ist die Ableitung $f'(x) = 2x$ und die Inverse ist $f^{-1}(x) = \sqrt{x}$. Die Ableitung der Inversen ist daher

$$(\sqrt{x})' = (f^{-1})'(x) = \frac{1}{f'(f^{-1}(x))} = \frac{1}{f'(\sqrt{x})} = \frac{1}{2\sqrt{x}}. \quad (5.18)$$

Beispiel: Gegeben $f(x) = \exp(x)$. Dann ist die Ableitung $f'(x) = \exp(x)$ und die Inverse ist $f^{-1}(x) = \ln(x)$. Die Ableitung der Inversen ist daher

$$(\ln x)' = (f^{-1})'(x) = \frac{1}{f'(f^{-1}(x))} = \frac{1}{\exp(\ln(x))} = \frac{1}{x}. \quad (5.19)$$

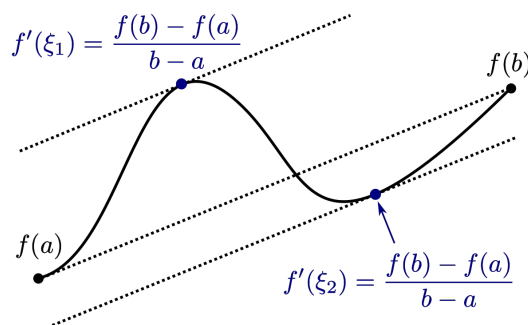


Abbildung 5.4: Mittelwertsatz der Differentialrechnung: Zwischen zwei Punkten eines Funktionsgraphen $(a, f(a))$ und $(b, f(b))$ einer differenzierbaren Funktion gibt es mindestens eine Stelle ξ , für welche die Tangente parallel zur Sekante durch die beiden Punkte ist, also die Ableitung exakt dem Vorwärts-Differenzenquotienten an der Stelle a mit $\Delta x = b - a$ entspricht. In der Abbildung gibt es sogar zwei solche Punkte ξ_1 und ξ_2 . Bildquelle: Wikipedia.

Mittelwertsatz und Regel von de L'Hospital

Ähnlich zum Zwischenwertsatz für stetige Funktionen gilt für differenzierbare Funktionen der

Mittelwertsatz der Differentialrechnung (ohne Beweis): Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig auf $[a, b]$ und differenzierbar auf $]a, b[$ mit $a < b$. Dann existiert (mindestens) ein $\xi \in]a, b[$, sodass

$$f'(\xi) = \frac{f(b) - f(a)}{b - a} \quad (5.20)$$

Figur 5.4 veranschaulicht den Mittelwertsatz. Ist eine Funktion bloß stetig, aber nicht differenzierbar, so gilt der Mittelwertsatz im allgemeinen nicht.

Aus einer Erweiterung des Mittelwertsatzes folgt die

Regel von de L'Hospital: Seien $f, g :]a, b[\rightarrow \mathbb{R}$ differenzierbar auf $]a, b[$. Es gebe ein $x_0 \in]a, b[$ mit entweder $f(x_0) = g(x_0) = 0$ oder $f(x_0) = \pm\infty$ oder $g(x_0) = \pm\infty$. Der Grenzwert des Quotienten $\frac{f(x)}{g(x)}$ ist dann unbestimmt mit $\frac{0}{0}$ oder $\pm\frac{\infty}{\infty}$. Es gelte $g'(x) \neq 0$ für alle $x \in]a, b[\setminus \{x_0\}$. Dann gilt

$$\lim_{x \rightarrow x_0} \frac{f(x)}{g'(x)} = c \Rightarrow \lim_{x \rightarrow x_0} \frac{f(x)}{g(x)} = c. \quad (5.21)$$

Hier kann c sowohl ein eigentlicher Grenzwert (dh. $c \in \mathbb{R}$) als auch ein uneigentlicher Grenzwert (dh. $c = \pm\infty$) sein.

Beweis: Der Fall $\lim_{x \rightarrow x_0} \frac{f(x)}{g(x)} = \pm\frac{\infty}{\infty}$ ist wegen $\frac{f(x)}{g(x)} = \frac{1/g(x)}{1/f(x)}$ auf den Fall $\lim_{x \rightarrow x_0} \frac{f(x)}{g(x)} = \frac{0}{0}$ reduzierbar. An der Stelle x_0 gilt dann $f(x_0) = g(x_0) = 0$. Wir nehmen an, dass sich x von links der Stelle x_0 nähert, also $x \leq x_0$. Wegen des Mittelwertsatzes existieren im

Intervall $]x, x_0[$ die Stellen ξ_f und ξ_g , sodass

$$\frac{f'(\xi_f)}{g'(\xi_g)} = \frac{f(x_0) - f(x)}{x_0 - x} \frac{x_0 - x}{g(x_0) - g(x)} = \frac{f(x_0) - f(x)}{g(x_0) - g(x)} = \frac{f(x)}{g(x)}. \quad (5.22)$$

Im Limes $x \rightarrow x_0$ gilt $\xi_f \rightarrow x_0$ und $\xi_g \rightarrow x_0$. Es folgt daher

$$\lim_{x \rightarrow x_0} \frac{f'(\xi_f)}{g'(\xi_g)} = \lim_{x \rightarrow x_0} \frac{f'(x)}{g'(x)} = \lim_{x \rightarrow x_0} \frac{f(x)}{g(x)}. \quad \square \quad (5.23)$$

Die Regel von de L'Hospital beruht auf dem Prinzip, dass bei einem zunächst unbestimmten Ausdruck der Grenzwert davon abhängt, ob die Funktion im Nenner oder die Funktion im Zähler die "stärker" wachsende bzw. fallende ist. Und diese Stärke wird durch die jeweilige Ableitung bestimmt.

Beispiel: Wir betrachten die Funktion $f(x) = \operatorname{sinc} x = \frac{\sin x}{x}$. Wir interessieren uns für den Grenzwert an der Stelle $x_0 = 0$. Wir erhalten zunächst den undefinierten Ausdruck

$$\lim_{x \rightarrow 0} f(x) = \lim_{x \rightarrow 0} \frac{\sin x}{x} = \frac{0}{0}. \quad (5.24)$$

Mit der Regel von de L'Hospital können wir den Grenzwert berechnen:

$$\lim_{x \rightarrow 0} f(x) = \lim_{x \rightarrow 0} \frac{\sin' x}{x'} = \lim_{x \rightarrow 0} \frac{\cos x}{1} = 1. \quad (5.25)$$

Beispiel: Wir betrachten die Funktion $f(x) = x \ln x$ und interessieren uns für den Grenzwert an der Stelle $x_0 = 0$. Zunächst erhalten wir den undefinierten Ausdruck

$$\lim_{x \rightarrow 0} f(x) = \lim_{x \rightarrow 0} \frac{\ln x}{1/x} = \frac{-\infty}{\pm\infty}. \quad (5.26)$$

Mit der Regel von de L'Hospital erhalten wir

$$\lim_{x \rightarrow 0} f(x) = \lim_{x \rightarrow 0} \frac{(\ln x)'}{(1/x)'} = \lim_{x \rightarrow 0} \frac{1/x}{-1/x^2} = \lim_{x \rightarrow 0} \frac{x}{-1} = 0. \quad (5.27)$$

Figur 5.5 veranschaulicht die letzten beiden Beispiele

Man beachte, dass die Regel von de L'Hospital nur angewendet werden darf, wenn der ursprüngliche Limes des Quotienten unbestimmt mit $\frac{0}{0}$ oder $\pm\frac{\infty}{\infty}$ ist. Wenn das nicht der Fall ist, führt die Regel im allgemeinen zu falschen Ergebnissen: Beispielsweise ist $\lim_{x \rightarrow 0} \frac{x^2+1}{x+1} = 1$. Aber $\lim_{x \rightarrow 0} \frac{(x^2+1)'}{(x+1)'} = \lim_{x \rightarrow 0} \frac{2x}{1} = 0$.

Newton-Verfahren

Das Newton-Verfahren ist eine numerische Methode zum Lösen von allgemeinen nicht-linearen Gleichungen und Gleichungssystemen. Im einfachsten Fall von nur einer Gleichung

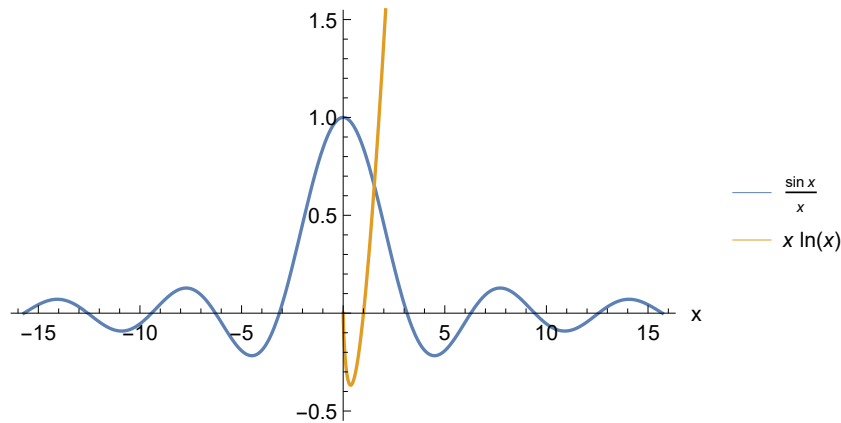


Abbildung 5.5: Die Funktionen $\text{sinc } x = \frac{\sin x}{x}$ (blau) und $x \ln x$ (orange) haben an der Stelle $x_0 = 0$ den Grenzwert 0, der sich mit der Regel von de L'Hospital bestimmen lässt.

chung und einer Variablen ist also die Aufgabe, die Lösungen von

$$f(x) = 0 \quad (5.28)$$

zu finden, also die Nullstellen einer gegebenen Funktion f .

Das Verfahren funktioniert wie folgt: Man startet bei einem (beliebigen) Ausgangspunkt x_1 . An dieser Stelle legt man die Tangente $t(x)$ an die Funktion:

$$t(x) = f(x_1) + f'(x_1)(x - x_1) \quad (5.29)$$

Nun berechnet man die Nullstelle x_2 dieser Tangente, setzt also $t(x_2) = f(x_1) + f'(x_1)(x_2 - x_1) = 0$. Die Lösung ist

$$x_2 = x_1 - \frac{f(x_1)}{f'(x_1)}. \quad (5.30)$$

Dies ist nun die nächst bessere Näherung für die Nullstelle der Funktion. Nun berechnet man die Tangente in x_2 und setzt das Verfahren Schritt für Schritt iterativ fort:

$$x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}. \quad (5.31)$$

Und zwar so lange, bis eine vorher festgesetzte Schranke $\varepsilon > 0$ erreicht ist:

$$|f(x_n)| < \varepsilon. \quad (5.32)$$

Das Newtonverfahren ist ein lokal konvergentes Verfahren, dh. es konvergiert, wenn der Startwert ausreichend Nahe bei der tatsächlichen Nullstelle liegt. In der Nähe der Nullstelle konvergiert das Verfahren quadratisch, dh. die Zahl der richtigen Dezimalstellen verdoppelt sich in jedem Schritt. Bei weit entferntem Startwert kann das Verfahren auch divergieren oder oszillieren, also die gleichen Funktionswerte immer wieder aufsuchen. Eine weitere

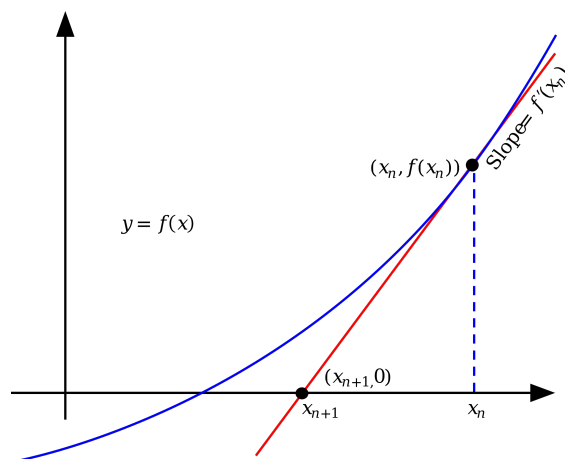


Abbildung 5.6: Newton-Verfahren: Gesucht ist die Nullstelle der (blauen) Funktion $f(x)$. Die Nullstelle x_{n+1} der (roten) Tangente mit Steigung (englisch: “slope”) $f'(x_n)$ an der Stelle x_n ist dafür eine bessere Näherung als x_n selbst. Bildquelle: Wikipedia.

Eigenschaft ist, dass das Verfahren pro Startwert nur eine Nullstelle finden kann.

Eine Veranschaulichung des Newton-Verfahrens finden Sie in Figur 5.6 sowie als Animation [hier](#).

Höhere Ableitungen und Extremstellen

Wenn die Ableitung einer Funktion $f(x)$, also $f'(x) = (f(x))'$, differenzierbar ist, dann können wir die zweite Ableitung $f''(x)$ bilden. Falls auch diese Funktion differenzierbar ist, dann existiert die dritte Ableitung $f'''(x)$. Und so weiter:

$$f'(x) = (f(x))', \quad (5.33)$$

$$f''(x) = (f'(x))', \quad (5.34)$$

$$f'''(x) = (f''(x))', \quad (5.35)$$

$$\vdots$$

$$f^{(n)}(x) = (f^{(n-1)}(x))'. \quad (5.36)$$

Hier wird die n -te Ableitung mit $f^{(n)}(x)$ bezeichnet. Diese Schreibweise mit Strichen und hochgestellter Nummerierung nennt man Lagrange-Notation.

Beispiel: Die Funktion $f(x) = x^2$ ist unendlich oft differenzierbar auf ganz $\mathbb{R}$. **Beweis:** Es gilt $f'(x) = 2x$, $f''(x) = 2$, $f'''(x) = f^{(4)} = f^{(5)} = \dots = 0$.

Beispiel: Die Funktion $f(x) = |x|^3$ ist zweimal differenzierbar auf ganz $\mathbb{R}$. Die dritte Ableitungsfunktion kann aber nicht mehr auf ganz $\mathbb{R}$ definiert werden. **Beweis:** In den Übungen.

Definition: Gegeben eine Funktion $f(x) : D \rightarrow \mathbb{R}$. Die Funktion f hat an der Stelle x_0

ein

- *lokales Minimum*, wenn es eine ε -Umgebung U_ε von x_0 gibt, sodass $f(x_0) \leq f(x)$ für alle $x \in U_\varepsilon$,
- *globales Minimum*, wenn $f(x_0) \leq f(x)$ für alle $x \in D$,
- *lokales Maximum*, wenn es eine ε -Umgebung U_ε von x_0 gibt, sodass $f(x_0) \geq f(x)$ für alle $x \in U_\varepsilon$,
- *globales Maximum*, wenn $f(x_0) \geq f(x)$ für alle $x \in D$.

Wir definieren ferner: Die Funktion f hat an der Stelle x_0 einen

- *Tiefpunkt = strenges lokales Minimum*, wenn es eine ε -Umgebung U_ε von x_0 gibt, sodass für alle $x \in U_\varepsilon$, $x \neq x_0$ gilt: $f(x_0) < f(x)$,
- *Hochpunkt = strenges lokales Maximum*, wenn es eine ε -Umgebung U_ε von x_0 gibt, sodass für alle $x \in U_\varepsilon$, $x \neq x_0$ gilt: $f(x_0) > f(x)$.

Minima bzw. Maxima bezeichnet man als Extremwerte. Die dazugehörigen Stellen x_0 nennt man Extremstellen. Die Punkte $(x_0, f(x_0))$, dh. (Extremstelle, Extremwert), nennt man Extrempunkte.

Differenzierbare Funktionen lassen sich auf Extremstellen untersuchen. Für Extremstellen gilt, dass dort die erste Ableitung verschwindet. Die zweite Ableitung entscheidet dann über die Art des Extrempunkts. Die Stelle x_0 ist ein

- Tiefpunkt, falls $f'(x_0) = 0$ und $f''(x_0) > 0$,
- Hochpunkt, falls $f'(x_0) = 0$ und $f''(x_0) < 0$.

Es ist aber auch möglich, dass trotz $f'(x_0) = 0$ kein Extrempunkt vorliegt. Kritische Stellen, dh. Stellen mit verschwindender erster Ableitung, die aber keine Extremstellen sind, sind Stellen von Sattelpunkten. Das ist zB. für die Funktion $f(x) = x^3$ an der Stelle $x_0 = 0$ der Fall, wo es einen sogenannten Sattelpunkt gibt. Die zweite Ableitung verschwindet, $f''(x_0) = 0$, und erst die dritte Ableitung $f'''(x_0) = 6$ ist ungleich 0. Die Funktion $f(x) = x^4$ hingegen hat an der Stelle $x_0 = 0$ trotz Verschwindens der ersten, zweiten und dritten Ableitung an der Stelle x_0 einen Tiefpunkt. Für x^5 liegt wieder ein Sattelpunkt vor, für x^6 ein Tiefpunkt (bzw. für $-x^6$ ein Hochpunkt). Und so weiter.

Es gilt allgemein: Sei f n -mal differenzierbar mit

$$f'(x_0) = f''(x_0) = \dots f^{(n-1)}(x_0) = 0, \quad f^{(n)} \neq 0 \quad (5.37)$$

Dann folgt

- für n gerade und $f^{(n)} > 0$: x_0 ist ein Tiefpunkt,
- für n gerade und $f^{(n)} < 0$: x_0 ist ein Hochpunkt,
- für n ungerade: x_0 ist ein Sattelpunkt (und damit kein Extrempunkt).

Beispiel: Wir betrachten die Funktion $f(x) = x^3 - 3x^2 - 1$. Die erste und zweite Ableitung sind

$$f'(x) = 3x^2 - 6x, \quad f''(x) = 6x - 6. \quad (5.38)$$

Wir setzen $f'(x) = 0$ und erhalten die beiden Lösungen $x_1 = 0$ und $x_2 = 2$. Die zweite

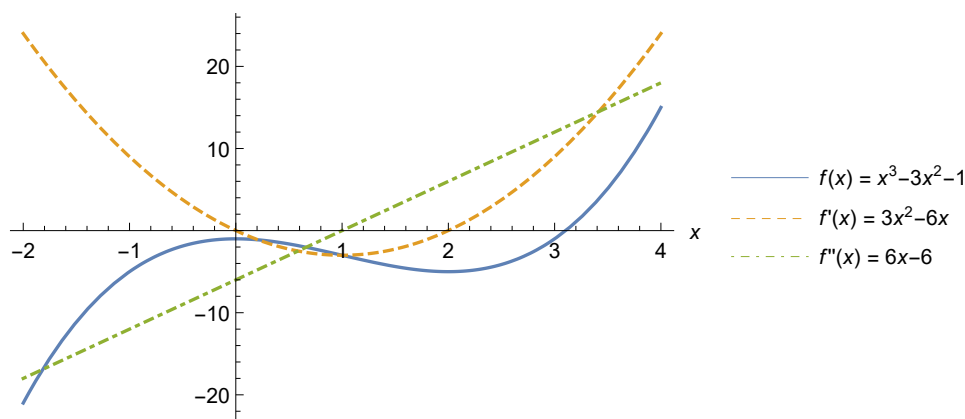


Abbildung 5.7: Die Funktion $f(x) = x^3 - 3x^2 - 1$ (blau) sowie ihre erste (orange) und zweite (grün) Ableitung. Die Funktion hat 2 Extremstellen: Einen Hochpunkt bei $x_1 = 0$ bzw. einen Tiefpunkt bei $x_2 = 2$. An diesen Stellen ist die erste Ableitung 0, und die zweite Ableitung ist negativ (bei x_1) bzw. positiv (bei x_2).

Ableitung ist an diesen Stellen $f''(x_1) = -6$ bzw. $f''(x_2) = 6$. Daher ist bei x_1 ein Hochpunkt und bei x_2 ein Tiefpunkt (siehe Figur 5.7).

Beispiel: Wir betrachten die Funktion $f(x) = (x-1)^5 = x^5 - 5x^4 + 10x^3 - 10x^2 + 5x - 1$. Die erste Ableitung $f'(x) = 5(x-1)^4$ hat nur eine Nullstelle, nämlich bei $x_0 = 1$. An dieser Stelle haben auch die zweite, dritte und vierte Ableitung eine Nullstelle: $f^{(2)}(x_0) = f^{(3)}(x_0) = f^{(4)}(x_0) = 0$. Erst $f^{(5)}(x_0) = 120$ ist ungleich 0. Da 5 ungerade ist, hat $f(x)$ bei $x_0 = 1$ einen Sattelpunkt. Die Funktion hat keine Extrempunkte.

Taylor-Reihen

Die *Taylor-Entwicklung* erlaubt es, mehrfach differenzierbare Funktionen durch Polynomfunktionen anzunähern. Für *glatte* (dh. unendlich oft differenzierbare) Funktionen führt dies zur *Taylor-Reihe*. Funktionen, die sich in jedem Punkt in eine Taylor-Reihe entwickeln lassen, heißen *analytisch*.

Die Intuition für die Taylor-Entwicklung ist folgende: Wir stellen uns eine beliebige mehrfach differenzierbare Funktion $f(x)$ vor, die wir an einer gewissen Stelle x_0 durch Polynome annähern wollen. In 0-ter Ordnung ist das die konstante Funktion

$$T_0(x, x_0) := f(x_0), \quad (5.39)$$

also die horizontale Gerade durch diesen Punkt. Das ist eine extrem grobe Näherung, und natürlich geht es besser, indem wir an der Stelle x_0 die Tangente an die Funktion anlegen, also eine Gerade durch den Punkt $(x_0, f(x_0))$ mit der Steigung $f'(x_0)$. Diese Gerade ist durch die lineare Funktion

$$T_1(x, x_0) := f(x_0) + f'(x_0)(x - x_0) \quad (5.40)$$

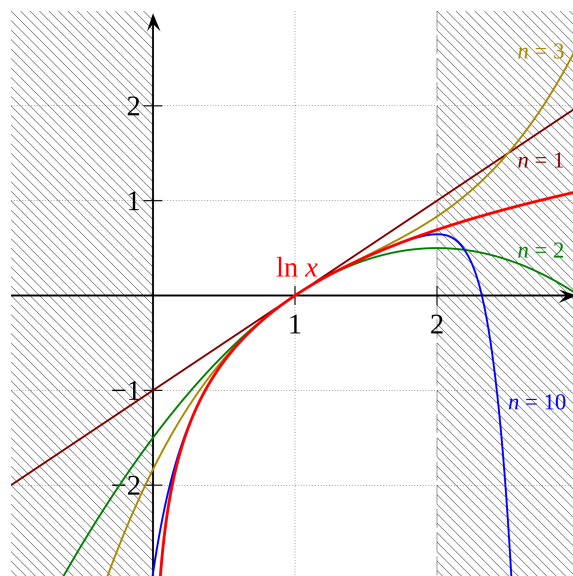


Abbildung 5.8: Taylor-Entwicklung der Funktion $\ln x$ (rot) an der Entwicklungsstelle $x_0 = 1$ durch Taylor-Polynome der Grade 1, 2, 3 und 10. Die Taylor-Entwicklung funktioniert hier nur im Intervall $x \in]0, 2]$. Bildquelle: Wikipedia.

gegeben. Dies ist das Taylor-Polynom erster Ordnung. Und es gilt natürlich $T_1'(x, x_0) = f'(x_0)$, wobei $'$ wie immer die Ableitung nach x bezeichnet.

In zweiter Ordnung wollen wir der Krümmung der Funktion Rechnung tragen. Wir konstruieren daher eine Schmiegeparabel, gegeben durch die quadratische Funktion (dh. Polynom zweiten Grades)

$$T_2(x, x_0) := f(x_0) + f'(x_0)(x - x_0) + \frac{f''(x_0)}{2}(x - x_0)^2, \quad (5.41)$$

wobei der Faktor 2 im Nenner sicherstellt, dass nun auch die zweite Ableitung von $T_2(x, x_0)$ genau mit der zweiten Ableitung von f an der Stelle x_0 übereinstimmt: $T_2''(x, x_0) = f''(x_0)$.

Wir können die Taylor-Entwicklung zu immer höheren Polynomgraden fortführen solange die höheren Ableitungen existieren.

Beispiel: Figur 5.8 zeigt die Taylor-Entwicklung der Funktion $\ln x$ an der Stelle $x_0 = 1$ mithilfe der Taylor-Polynome $T_1(x, x_0)$, $T_2(x, x_0)$, $T_3(x, x_0)$ und $T_{10}(x, x_0)$. Eine Animation für $f(x) = \ln(1 + x)$ um die Entwicklungsstelle $x_0 = 0$ findet man *hier* (mit $N \equiv n$).

Beispiel: Eine Animation für die Entwicklung der Exponentialfunktion $\exp(x)$ an der Stelle $x_0 = 0$ finden Sie *hier*.

Der Satz von Taylor formalisiert die obigen Überlegungen, wenn wir bis zum Grad n entwickeln. Er gibt uns zudem eine Abschätzung für den Approximationsfehler, also für die Differenz des tatsächlichen Funktionswerts $f(x)$ vom Taylor-Polynom $T_n(x, x_0)$.

Satz von Taylor: Sei $f :]a, b[\rightarrow \mathbb{R}$ eine $(n+1)$ -mal differenzierbare Funktion auf $[a, b]$. Dann kann man die Funktion an jeder Stelle $x \in]a, b[$ durch Entwicklung um jede Stelle $x_0 \in]a, b[$ wie folgt darstellen:

$$f(x) = f(x_0) + \frac{f'(x_0)}{1!} (x - x_0) + \frac{f''(x_0)}{2!} (x - x_0)^2 + \dots + \frac{f^{(n)}(x_0)}{n!} (x - x_0)^n + R_{n+1}(x, x_0) \quad (5.42)$$

mit

$$R_{n+1}(x, x_0) = \frac{f^{(n+1)}(\xi)}{(n+1)!} (x - x_0)^{n+1}, \quad (5.43)$$

wobei ξ zwischen x und x_0 liegt. Die ersten $n+1$ Terme sind das Taylor-Polynom $T_n(x, x_0)$, und R_{n+1} ist das Restglied, das eine Fehlerabschätzung der Approximation durch das Taylor-Polynom erlaubt.

Für unendlich oft differenzierbare Funktionen $f(x)$ mit $\lim_{n \rightarrow \infty} R_{n+1}(x, x_0) = 0$ bekommt man die Darstellung durch die (unendliche) Taylor-Reihe:

$$\begin{aligned} f(x) &= f(x_0) + \frac{f'(x_0)}{1!} (x - x_0) + \frac{f''(x_0)}{2!} (x - x_0)^2 + \dots \\ &= \sum_{n=0}^{\infty} \frac{f^{(n)}(x_0)}{n!} (x - x_0)^n. \end{aligned} \quad (5.44)$$

Diese Reihenentwicklung ermöglicht es insbesondere, die Koeffizienten in der Potenzreihendarstellung einer Funktion zu finden.

Beispiel: Wir nehmen an, dass wir die Potenzreihendarstellung der Exponentialfunktion $f(x) = \exp(x)$ noch nicht kennen. Aber wir wissen, dass $\exp'(x) = \exp(x)$ und $\exp(0) = 1$ gilt. Dann können wir an der Stelle $x_0 = 0$ die Taylor-Reihe^a bilden:

$$f(x) = \sum_{n=0}^{\infty} \frac{f^{(n)}(x_0)}{n!} (x - x_0)^n = \sum_{n=0}^{\infty} \frac{\exp^{(n)}(0)}{n!} x^n = \sum_{n=0}^{\infty} \frac{1}{n!} x^n. \quad (5.45)$$

Wir haben damit die Potenzreihendarstellung von $\exp(x)$ hergeleitet. Betrachten wir diese Taylor-Entwicklung beispielsweise nur bis zum zweiten Taylor-Polynom, so bekommen wir eine Approximation der Exponentialfunktion wie folgt:

$$\exp(x) = 1 + x + \frac{x^2}{2} + \mathcal{O}(x^3). \quad (5.46)$$

Man nennt das auch eine Näherung in zweiter Ordnung. Das Landau-Symbol (auch genannt: O-Notation) $\mathcal{O}(g)$ bezeichnet abstrakt einen Term, der höchstens genau so schnell wächst wie g . In unserem Fall ist dieser Term konkret das Restglied R_3 .

^aIm Spezialfall $x_0 = 0$ nennt man die Taylor-Reihe auch Maclaurin-Reihe.

Beispiel: Wir suchen die Taylor-Entwicklung von $f(x) = \cos(x)$ an der Stelle $x_0 = 0$ bis zur 3. Ordnung. *Lösung:*

$$\begin{aligned}\cos(x) &= \cos(x_0) + \cos'(x_0)(x-x_0) + \frac{\cos''(x_0)}{2}(x-x_0)^2 + \frac{\cos'''(x_0)}{3!}(x-x_0)^3 + \mathcal{O}(x^4) \\ &= 1 - \sin(0)x - \frac{\cos(0)}{2}x^2 + \frac{\sin(0)}{3!}x^3 + \mathcal{O}(x^4) \\ &= 1 - \frac{x^2}{2} + \mathcal{O}(x^4).\end{aligned}\tag{5.47}$$

Beispiel: Leiten Sie die Potenzreihendarstellung der Funktion $f(x) = \sin x$ durch Taylor-Entwicklung an der Stelle $x_0 = 0$ her. *Lösung:* In den Übungen.

Leibniz-Notation

Oftmals deutlich intuitiver und praktischer als die Lagrange-Notation ist die Leibniz-Notation. Leibniz führte dafür die *Differentiale* $dy = df(x)$ bzw. dx ein, die symbolisch für eine infinitesimal kleine Änderung der Funktion $y = f(x)$ bzw. des Arguments x stehen. Die erste Ableitung einer Funktion ist der Grenzwert des Differenzenquotienten, also:

$$f'(x) = y' = \lim_{\Delta x \rightarrow 0} \frac{\Delta y}{\Delta x} = \frac{dy}{dx} = \frac{df(x)}{dx}\tag{5.48}$$

Die letzten beiden Ausdrücke sind die Leibniz-Notation und werden gesprochen als “d y nach d x” bzw. “d f von x nach d x”.

Der *Ableitungsoperator* $\frac{d}{dx}$ kann nun verwendet werden, um nicht nur die erste sondern auch höhere Ableitungen zu bilden:

$$\frac{d}{dx}y = \frac{dy}{dx},\tag{5.49}$$

$$\left(\frac{d}{dx}\right)^2 y = \frac{d}{dx} \frac{dy}{dx} = \frac{d^2 y}{dx^2},\tag{5.50}$$

$$\left(\frac{d}{dx}\right)^3 y = \frac{d}{dx} \frac{d^2 y}{dx^2} = \frac{d^3 y}{dx^3},\tag{5.51}$$

$\vdots$

$$\left(\frac{d}{dx}\right)^n y = \frac{d}{dx} \frac{d^{n-1} y}{dx^{n-1}} = \frac{d^n y}{dx^n}.\tag{5.52}$$

In der Leibniz-Notation nehmen die Ketten- und Umkehrregel eine ausgesprochen intuitive Form an. Sei $y = y(u)$ eine Funktion von u und $u = u(x)$ eine Funktion von x . Dann gilt die Kettenregel

$$\frac{dy}{dx} = \frac{dy}{du} \frac{du}{dx}.\tag{5.53}$$

Die Umkehrfunktion von $y(x)$ ist $x(y)$. Die Ableitung der Umkehrfunktion erhält man

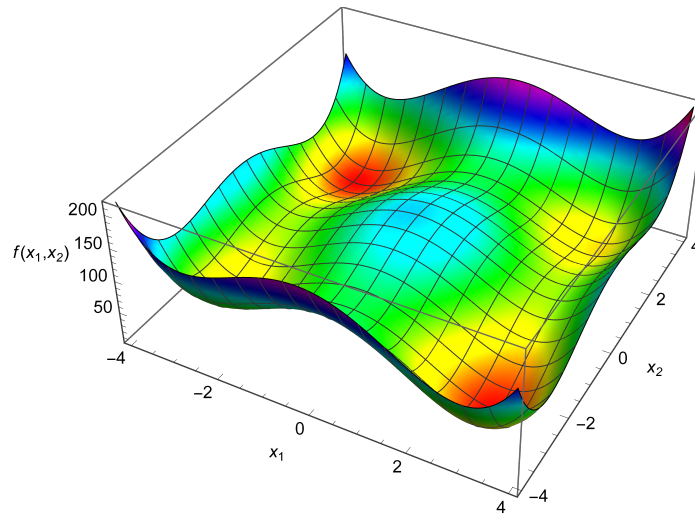


Abbildung 5.9: Die Polynom-Funktion (5.56) hat 9 kritische Punkte: Sie hat 4 lokale Minima (gelb bzw. rot), zwei davon sind global (rot). Sie hat ein lokales Maximum in der Mitte (hellblau). Und sie hat 4 Sattelpunkte (grün), die zwischen den Minima liegen.

durch die Umkehrregel:

$$\frac{dx}{dy} = \frac{1}{\frac{dy}{dx}}. \quad (5.54)$$

5.2 In mehreren Variablen

Wir betrachten nun reellwertige n -dimensionale Funktionen f , die n Variablen $x_1, x_2, \dots, x_n$ auf eine reelle Zahl abbilden:

$$f: \mathbb{R}^n \rightarrow \mathbb{R}, (x_1, x_2, \dots, x_n) \mapsto f(x_1, x_2, \dots, x_n). \quad (5.55)$$

Beispiel: Die Funktion

$$f(x_1, x_2) = x_1^4 + x_2^4 - 16x_1^2 - 12x_2^2 + 2x_1x_2 + 120 \quad (5.56)$$

ist ein zweidimensionales (auch genannt: bivariates) Polynom vierter Ordnung (siehe Figur 5.9).

Im folgenden verwenden wir oft die Kurzschreibweise $\mathbf{x} \equiv (x_1, x_2, \dots, x_n)$. Wir können eine n -dimensionale Funktion entlang ihrer n einzelnen Variablen (Koordinatenachsen) ableiten. Dazu betrachtet man jeweils alle anderen Variablen als konstant. Diese “teilweisen” Differentiationen bezeichnet man daher als *partielle Ableitungen* und verwendet das

alternative Symbol ∂ für die Differentiale:¹

$$\frac{\partial f(\mathbf{x})}{\partial x_1} = \lim_{h \rightarrow 0} \frac{f(x_1 + h, x_2, x_3, \dots, x_{n-1}, x_n) - f(x_1, x_2, x_3, \dots, x_{n-1}, x_n)}{h}, \quad (5.57)$$

$$\frac{\partial f(\mathbf{x})}{\partial x_2} = \lim_{h \rightarrow 0} \frac{f(x_1, x_2 + h, x_3, \dots, x_{n-1}, x_n) - f(x_1, x_2, x_3, \dots, x_{n-1}, x_n)}{h}, \quad (5.58)$$

$\vdots$

$$\frac{\partial f(\mathbf{x})}{\partial x_n} = \lim_{h \rightarrow 0} \frac{f(x_1, x_2, x_3, \dots, x_{n-1}, x_n + h) - f(x_1, x_2, x_3, \dots, x_{n-1}, x_n)}{h}. \quad (5.59)$$

Wenn die partielle Ableitung entlang von x_i existiert, mit $i = 1, 2, \dots, n$, dann nennt man die Funktion partiell differenzierbar in x_i .

Wenn wir die Einheitsvektoren $\mathbf{e}_i = (0, 0, \dots, 0, 1_i, 0, 0, \dots, 0)$ entlang der i -ten Koordinatenachse definieren, mit $i = 1, 2, \dots, n$, dann lassen sich die partiellen Ableitungen recht kompakt schreiben als

$$\frac{\partial f(\mathbf{x})}{\partial x_i} = \lim_{h \rightarrow 0} \frac{f(\mathbf{x} + h\mathbf{e}_i) - f(\mathbf{x})}{h}. \quad (5.60)$$

Gradient und Hesse-Matrix

Der Vektor aller partiellen Ableitungen einer Funktion f ist der *Gradient* der Funktion, geschrieben als $\text{grad}(f)$. Eine weit verbreitete Schreibweise ist gegeben durch den *Nabla-Operator*

$$\nabla \equiv \frac{\partial}{\partial \mathbf{x}} \equiv \left(\frac{\partial}{\partial x_1}, \frac{\partial}{\partial x_2}, \dots, \frac{\partial}{\partial x_n} \right)^T, \quad (5.61)$$

dh. den Vektor aller partiellen Ableitungsoperatoren. Wir verwenden die in den Ingenieur- und Naturwissenschaften übliche Konvention von Spaltenvektoren. Und um Platz zu sparen, schreiben wir oft Zeilenvektoren, die wir am Ende mit dem hochgestellten Symbol T transponieren. Dies führt uns zu folgender

Definition: Der Gradient einer Funktion $f(\mathbf{x})$ ist gegeben durch

$$\text{grad} f(\mathbf{x}) \equiv \nabla f(\mathbf{x}) \equiv \frac{\partial f(\mathbf{x})}{\partial \mathbf{x}} = \left(\frac{\partial f(\mathbf{x})}{\partial x_1}, \frac{\partial f(\mathbf{x})}{\partial x_2}, \dots, \frac{\partial f(\mathbf{x})}{\partial x_n} \right)^T. \quad (5.62)$$

Anschaulich ist der Gradient einer Funktion an der Stelle $\mathbf{x} = (x_1, x_2, \dots, x_n)$ jener Vektor, der an dieser Stelle in die Richtung des steilsten Anstiegs der Funktion zeigt und dessen Länge die Stärke der Steigung misst.

Wenn die partiellen Ableitungen selbst wieder differenzierbar sind, können wir die zweiten Ableitungen bilden. Diese werden in der sogenannten Hesse-Matrix zusammengefasst.

¹Wir nehmen hier den Limes der Schrittweite h gegen null, und nicht mehr den Limes des Index n der Nullfolgen h_n gegen unendlich. Die Variable n benötigen wir nun für die Dimension der Funktion.

Definition: Die Hesse-Matrix einer Funktion $f(\mathbf{x})$ ist gegeben durch

$$H_f(\mathbf{x}) = \begin{pmatrix} \frac{\partial^2 f(\mathbf{x})}{\partial x_1^2} & \frac{\partial^2 f(\mathbf{x})}{\partial x_1 \partial x_2} & \cdots & \frac{\partial^2 f(\mathbf{x})}{\partial x_1 \partial x_n} \\ \frac{\partial^2 f(\mathbf{x})}{\partial x_2 \partial x_1} & \frac{\partial^2 f(\mathbf{x})}{\partial x_2^2} & \cdots & \frac{\partial^2 f(\mathbf{x})}{\partial x_2 \partial x_n} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2 f(\mathbf{x})}{\partial x_n \partial x_1} & \frac{\partial^2 f(\mathbf{x})}{\partial x_n \partial x_2} & \cdots & \frac{\partial^2 f(\mathbf{x})}{\partial x_n^2} \end{pmatrix}. \quad (5.63)$$

Hier bedeutet $\frac{\partial^2 f(\mathbf{x})}{\partial x_i \partial x_j} = \frac{\partial}{\partial x_i} \frac{\partial}{\partial x_j} f(\mathbf{x})$, dass zuerst nach x_j abgeleitet wird, danach nach x_i .

Die kritischen Stellen der Funktion $f(\mathbf{x})$ bekommen wir, indem wir *alle* ersten partiellen Ableitungen *zugleich* auf null setzen. An jeder dieser Stellen berechnet man die Hesse-Matrix. Ist die Matrix dort positiv definit, so befindet sich an dieser Stelle ein Tiefpunkt. Ist sie negativ definit, so ist dort ein Hochpunkt. Ist sie indefinit, dann hat die Funktion an der Stelle einen Sattelpunkt. Bei positiver Semi-Definitheit liegt entweder ein Minimum oder ein Sattelpunkt vor. Bei negativer Semi-Definitheit liegt entweder ein Maximum oder ein Sattelpunkt vor.

Beispiel: Berechne die Extrempunkte der Funktion (5.56). **Lösung:** Die Funktion ist in beiden Variablen partiell differenzierbar. Der Gradient ist

$$\nabla f(x_1, x_2) = \begin{pmatrix} \frac{\partial f(x_1, x_2)}{\partial x_1} \\ \frac{\partial f(x_1, x_2)}{\partial x_2} \end{pmatrix} = \begin{pmatrix} 4x_1^3 - 32x_1 + 2x_2 \\ 4x_2^3 - 24x_2 + 2x_1 \end{pmatrix}. \quad (5.64)$$

Die partiellen Ableitungen sind wiederum alle differenzierbar. Die Hesse-Matrix hat die Form^a

$$H_f(x_1, x_2) = \begin{pmatrix} \frac{\partial^2 f(x_1, x_2)}{\partial x_1^2} & \frac{\partial^2 f(x_1, x_2)}{\partial x_1 \partial x_2} \\ \frac{\partial^2 f(x_1, x_2)}{\partial x_2 \partial x_1} & \frac{\partial^2 f(x_1, x_2)}{\partial x_2^2} \end{pmatrix} = \begin{pmatrix} 12x_1^2 - 32 & 2 \\ 2 & 12x_2^2 - 24 \end{pmatrix}. \quad (5.65)$$

Null-Setzen von (5.64) liefert zwei kubische Gleichungen mit zwei Variablen. Die 9 Lösungen (berechnet mit Computer-Hilfe) sind unsere kritischen Stellen, dh. Extremstellen und Sattelpunkte:

$$x_1 = 0 \quad x_2 = 0 \quad f(x_1, x_2) = 120 \quad (\text{H}) \quad (5.66)$$

$$x_1 = -0.153 \quad x_2 = -2.443 \quad f(x_1, x_2) = 84.37 \quad (\text{S}) \quad (5.67)$$

$$x_1 = +0.153 \quad x_2 = +2.443 \quad f(x_1, x_2) = 84.37 \quad (\text{S}) \quad (5.68)$$

$$x_1 = -2.753 \quad x_2 = -2.326 \quad f(x_1, x_2) = 33.33 \quad (\text{T}) \quad (5.69)$$

$$x_1 = +2.753 \quad x_2 = +2.326 \quad f(x_1, x_2) = 33.33 \quad (\text{T}) \quad (5.70)$$

$$x_1 = -2.821 \quad x_2 = -0.237 \quad f(x_1, x_2) = 56.67 \quad (\text{S}) \quad (5.71)$$

$$x_1 = +2.821 \quad x_2 = +0.237 \quad f(x_1, x_2) = 56.67 \quad (\text{S}) \quad (5.72)$$

$$x_1 = -2.905 \quad x_2 = +2.563 \quad f(x_1, x_2) = 5.63 \quad (\text{T}) \quad (5.73)$$

$$x_1 = +2.905 \quad x_2 = -2.563 \quad f(x_1, x_2) = 5.63 \quad (\text{T}) \quad (5.74)$$

Hier haben wir schon Hochpunkte mit (H), Sattelpunkte mit (S) und Tiefpunkte mit (T) gekennzeichnet, müssen diese Eigenschaften aber noch beweisen. Wir machen dies repräsentativ für drei dieser Stellen, nämlich die erste $\mathbf{x}^{(1)} = (0, 0)$, zweite $\mathbf{x}^{(2)} = (-0.153, +2.443)$ und neunte $\mathbf{x}^{(9)} = (+2.905, -2.563)$. Die Hesse-Matrizen an diesen Stellen sind:

$$H^{(1)} = H_f(0, 0) = \begin{pmatrix} -32 & 2 \\ 2 & -24 \end{pmatrix}, \quad (5.75)$$

$$H^{(2)} = H_f(-0.153, +2.443) = \begin{pmatrix} -31.72 & 2 \\ 2 & 47.62 \end{pmatrix}, \quad (5.76)$$

$$H^{(9)} = H_f(+2.905, -2.563) = \begin{pmatrix} 69.29 & 2 \\ 2 & 54.80 \end{pmatrix}. \quad (5.77)$$

Eine symmetrische (2×2) -Matrix $A = \begin{pmatrix} a & b \\ b & d \end{pmatrix}$ hat die beiden führenden Hauptminoren $A_1 = a$ und $A_2 = |A| = \det(A) = ad - b^2$. Die Matrix ist positiv definit, genau dann wenn $A_1 > 0$ und $A_2 > 0$. Die Matrix ist negativ definit, genau dann wenn $A_1 < 0$ und $A_2 > 0$. Sie ist indefinit, falls keiner der beiden vorigen Fälle zutrifft und wenn sie ferner nicht semi-definit ist, also wenn ferner $A_2 \neq 0$. In unseren drei Fällen haben wir:

$$H^{(1)}: A_1 = -32 < 0, A_2 = 764 > 0 \Rightarrow \text{negativ definit} \Rightarrow \text{Hochpunkt},$$

$$H^{(2)}: A_1 = -31.72 < 0, A_2 = -1514.6 < 0 \Rightarrow \text{indefinit} \Rightarrow \text{Sattelpunkt},$$

$$H^{(9)}: A_1 = 69.29 > 0, A_2 = 3793.4 > 0 \Rightarrow \text{positiv definit} \Rightarrow \text{Tiefpunkt}.$$

^aEs ist kein Zufall, dass $\frac{\partial^2 f(x_1, x_2)}{\partial x_1 \partial x_2} = \frac{\partial^2 f(x_2, x_1)}{\partial x_2 \partial x_1}$. Für zweimal differenzierbare Funktionen mit stetigen zweiten Ableitungen garantiert der *Satz von Schwarz*, dass die Reihenfolge der partiellen Ableitungen vertauscht werden kann.

Gradientenabstiegsverfahren

Das Verfahren des steilsten Abstiegs (englisch: “gradient descent”) ist eine Methode der numerischen Mathematik, um insbesondere multidimensionale Funktionen zu minimieren. Dabei werden, ausgehend von einem Startpunkt, kleine Schritte in Richtung des negativen Gradienten – also in Richtung des steilsten Abstiegs der Funktion – gemacht. Bildlich gesprochen, ähnelt das Verfahren einer Bergsteigerin, die in starkem Nebel möglichst schnell ins Tal möchte. Da sie nur ihre nächste Umgebung kennt, setzt Sie ihren nächsten Schritt stets in Richtung des steilsten Abstiegs. Sie wird früher oder später ein (lokales) Minimum ihrer Landschaft finden.

Der Gradientenabstieg wurde im Jahr 1847 von Cauchy vorgeschlagen und hat heutzutage zahllose Anwendungsgebiete. Insbesondere bildet er das Rückgrat des Trainingsprozesses von künstlichen neuronalen Netzwerken. Dort muss eine n -dimensionale Kosten- bzw. Fehlerfunktion minimiert werden, wobei n die Zahl der Parameter ist. In modernen Large Language Models kann n in der Größenordnung von 100 Milliarden = 10^{11} liegen.

Formal funktioniert das Verfahren wie folgt: Gegeben eine Funktion f , von der man ein Minimum finden möchte. Wir starten an einer beliebigen Stelle $\mathbf{x}_0$ und berechnen dort – beispielsweise mittels Differenzenquotienten – den Gradienten $\nabla f(\mathbf{x}_0) = \frac{\partial f(\mathbf{x}_0)}{\partial \mathbf{x}}$. Nun geht man mit der Schrittweite $\lambda > 0$ einen Schritt in die Richtung des negativen Gradienten und landet bei der neuen Stelle $\mathbf{x}_1$. Dies wird iterativ wiederholt. Die allgemeine Update-Regel lautet also:

$$\mathbf{x}_i = \mathbf{x}_{i-1} - \lambda \frac{\partial f(\mathbf{x}_{i-1})}{\partial \mathbf{x}}, \quad (5.78)$$

mit $i \in \mathbb{N}$. Dies wird solange wiederholt, bis ein Abbruchkriterium erreicht ist, beispielsweise eine bestimmte Anzahl $i_{\max}$ von Schritten (dh. $i = i_{\max}$) oder die Tatsache, dass sich der Funktionswert von einem zum nächsten Schritt weniger ändert als eine vorgegebene Schranke $\varepsilon > 0$ (dh. $|f(\mathbf{x}_i) - f(\mathbf{x}_{i-1})| < \varepsilon$).

Die richtige Wahl der Schrittweite ist von Bedeutung. Bei zu kleiner Schrittweite dauert die Optimierung zu lange. Bei zu großer wird das Minimum immer wieder verfehlt.

6. Integralrechnung

Die Integralrechnung und Differentialrechnung bilden die beiden Säulen der Infinitesimalrechnung – also der mathematischen Untersuchung kontinuierlicher Veränderungen. Das bestimmte Integral einer Funktion gibt den Flächeninhalt an, den die Funktion mit der x -Achse einschließt. Das unbestimmte Integral einer Funktion f ist eine Funktion F , deren erste Ableitung f ist. Das Integrieren ist also die Umkehrung des Differenzierens.

Im Gegensatz zur Differentiation gibt es für die Integration auch elementarer Funktionen oftmals keine einfachen Algorithmen – es gilt der Spruch “Differenzieren ist ein Handwerk, Integrieren eine Kunst”.

6.1 In einer Variablen

Definition: Gegeben eine Funktion $f :]a, b[\rightarrow \mathbb{R}$. Die differenzierbare Funktion $F :]a, b[\rightarrow \mathbb{R}$ heißt *Stammfunktion* von f , wenn für alle $x \in]a, b[$ gilt:

$$\frac{d}{dx}F(x) = f(x). \quad (6.1)$$

Falls $F(x)$ eine Stammfunktion von $f(x)$ ist, dann ist auch $F(x) + c$ mit $c \in \mathbb{R}$ eine Stammfunktion, da beim Differenzieren die Konstante c wegfällt. Für die Gesamtheit aller Stammfunktionen schreibt man das *unbestimmte Integral* der Funktion $f(x)$:

$$\int f(x) dx = F(x) + c. \quad (6.2)$$

Beispiele: Wir fassen einige Funktionen und deren Stammfunktionen tabellarisch zusammen:

$f(x)$	$\int f(x) dx$
0	c
x^n	$\frac{x^{n+1}}{n+1} + c, n \in \mathbb{R} \setminus \{-1\}$
$\frac{1}{x}$	$\ln(x) + c$
e^x	$e^x + c$
$\ln(x)$	$x(\ln(x) - 1) + c$
$\sin(x)$	$-\cos(x) + c$
$\cos(x)$	$\sin(x) + c$

Integrationsregeln

Es gelten folgende Integrationsregeln:

- Summenregel: $\int [f(x) + g(x)] dx = \int f(x) dx + \int g(x) dx$
- Faktorregel: $\int \lambda f(x) dx = \lambda \int f(x) dx$
- Partielle Integration: $\int f'(x) g(x) dx = f(x) g(x) - \int f(x) g'(x) dx$
- Integration durch Substitution: $\int f(g(x)) g'(x) dx = \int f(u) du|_{u \rightarrow g(x)}$

Beweis der partiellen Integration:

$$\begin{aligned}
 \int f(x) g'(x) dx + \int f'(x) g(x) dx &= \int [f(x) g'(x) + f'(x) g(x)] dx \\
 &= \int [f(x) g(x)]' dx \\
 &= f(x) g(x) + c.
 \end{aligned} \tag{6.3}$$

Beweis der Integration durch Substitution: Mit $u := g(x)$ gilt $\frac{du}{dx} = g'(x)$ und daher $du = g'(x) dx$. Daraus folgt

$$\int f(g(x)) g'(x) dx = \int f(u) du. \tag{6.4}$$

Beispiel: Berechne $\int x e^x dx$. *Lösung:* Durch partielle Integration mit $f'(x) = e^x$ und $g(x) = x$. Es gilt $f(x) = e^x$ und $g'(x) = 1$. Damit folgt

$$\int x e^x dx = x e^x - \int e^x dx = x e^x - e^x + c = e^x (x - 1) + c. \tag{6.5}$$

Beispiel: Berechne $\int x \cos(x^2 + 1) dx$. *Lösung:* Durch Substitution $u := x^2 + 1$ gilt $du = 2x dx$. Damit folgt

$$\int x \cos(x^2 + 1) dx = \frac{1}{2} \int \cos(u) du = \frac{1}{2} \sin(u) + c = \frac{1}{2} \sin(x^2 + 1) + c. \tag{6.6}$$

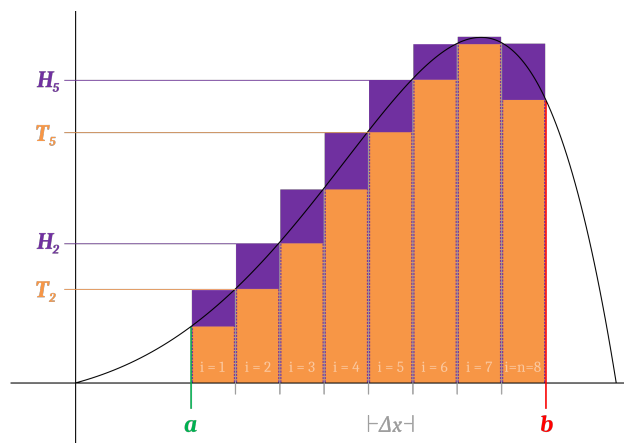


Abbildung 6.1: Die Fläche, die eine Funktion im Intervall $[a, b]$ mit der x -Achse einschließt, kann durch eine Treppenfunktionen angenähert werden, die auf jedem Teilintervall mit Index $i = 1, \dots, n$ konstant das Supremum H_i (violett) bzw. das Infimum T_i (orange) der Funktion ist. Bildquelle: Wikipedia.

Riemann-Integration

Wir wenden uns nun dem Problem zu, den Flächeninhalt unter einer Funktion zu berechnen, also die Fläche zwischen der x -Achse und einer Funktion $f(x)$ in einem Intervall $[a, b]$. Die Methode der Ober- und Untersummen – vorgeschlagen von Darboux und Riemann – besteht darin, für das Intervall eine Zerlegung Z in n Teile mit den Stützstellen $a = x_0 < x_1 < x_2 < \dots < x_n = b$ vorzunehmen. In jedem Teilintervall $I_i := [x_{i-1}, x_i]$ mit Index $i = 1, \dots, n$ kann man nun die Fläche unter der Funktion durch ein Rechteck annähern, wobei man als Höhe entweder das Supremum der Funktion im Teilintervall oder das Infimum nehmen kann. (In Figur 6.1 ist der Spezialfall dargestellt, in dem alle Teilintervalle die gleiche Länge Δx haben.) Der erste Fall führt zur Obersumme, der zweite zur Untersumme:

$$O(Z) := \sum_{i=1}^n (x_i - x_{i-1}) \sup_{x_{i-1} \leq x \leq x_i} f(x), \quad (6.7)$$

$$U(Z) := \sum_{i=1}^n (x_i - x_{i-1}) \inf_{x_{i-1} \leq x \leq x_i} f(x). \quad (6.8)$$

Bei einer Verfeinerung der Zerlegung Z , dh. Vergrößerung von n , wird die Obersumme kleiner und die Untersumme größer. Wir betrachten nun alle möglichen Zerlegungen und definieren das obere bzw. untere Darbouxsche Integral von f :

$$\overline{\int_a^b} f(x) dx := \inf_Z O(Z), \quad (6.9)$$

$$\underline{\int_a^b} f(x) dx := \sup_Z U(Z). \quad (6.10)$$

Definition: Falls das obere und untere Darboux'sche Integral gleich sind, dann heißt f *Riemann-integrierbar* (oder Darboux-integrierbar) und der gemeinsame Wert ist das *bestimmte Integral* – genannt *Riemannsches Integral* (oder Darboux'sches Integral) – von f über dem Intervall $[a, b]$:

$$\int_a^b f(x) dx := \overline{\int_a^b f(x) dx} = \underline{\int_a^b f(x) dx}. \quad (6.11)$$

Die reellen Zahlen a und b heißen Integrationsgrenzen, die zu integrierende Funktion $f(x)$ heißt Integrand, die Variable x heißt Integrationsvariable, und dx ist das Differential, das symbolisch einen infinitesimal kleinen Intervall darstellt.

Symbolisch kann man den Übergang von der Riemann-Summe zum Riemann-Integral also wie folgt schreiben:

$$\sum_{i=1}^n f(x_i) \Delta x_i \xrightarrow[\Delta x_i \rightarrow 0]{n \rightarrow \infty} \int_a^b f(x) dx. \quad (6.12)$$

Beispiel: Berechne das Integral $\int_0^1 x^2 dx$ mittels Ober- bzw. Untersumme. *Lösung:* Wir wählen die Zerlegung Z des Integrationsintervalls in n gleiche Teile der Länge $\frac{1}{n}$. Die Funktion x^2 ist monoton steigend: Der größte Funktionswert im i -ten Teilintervall ist am Intervall-Ende bei $(\frac{i}{n})^2$, und der kleinste ist am Intervall-Anfang bei $(\frac{i-1}{n})^2$. Also:

$$O(Z) = \sum_{i=1}^n \frac{1}{n} \left(\frac{i}{n}\right)^2 = \frac{1}{n^3} \sum_{i=1}^n i^2 = \frac{1}{n^3} \frac{n(n+1)(2n+1)}{6} = \frac{1}{3} + \frac{1}{2n} + \frac{1}{6n^2}, \quad (6.13)$$

$$U(Z) = \sum_{i=1}^n \frac{1}{n} \left(\frac{i-1}{n}\right)^2 = \frac{1}{n^3} \sum_{i=0}^{n-1} i^2 = \frac{1}{n^3} \frac{(n-1)n(2n-1)}{6} = \frac{1}{3} - \frac{1}{2n} + \frac{1}{6n^2}. \quad (6.14)$$

Für unendlich feine Zerlegung konvergieren Ober- und Untersumme gegen den gleichen Grenzwert. Mithin ist das Riemannsche Integral gegeben durch

$$\int_0^1 x^2 dx = \lim_{n \rightarrow \infty} O(Z) = \lim_{n \rightarrow \infty} U(Z) = \frac{1}{3}. \quad (6.15)$$

Im folgenden Kapitel werden wir sehen, wie wir bestimmte Integrale auch ohne den Grenzwert von Riemann-Summen berechnen können.

Hauptsatz der Differential- und Integralrechnung

Der Hauptsatz der Differential- und Integralrechnung (auch genannt: Fundamentalsatz der Analysis) verbindet die beiden Konzepte der Differentiation und der Integration. Der Hauptsatz besteht aus zwei Teilen, die mitunter auch erster und zweiter Hauptsatz genannt werden.

Erster Teil bzw. erster Hauptsatz: Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig. Dann ist für jedes $c \in [a, b]$

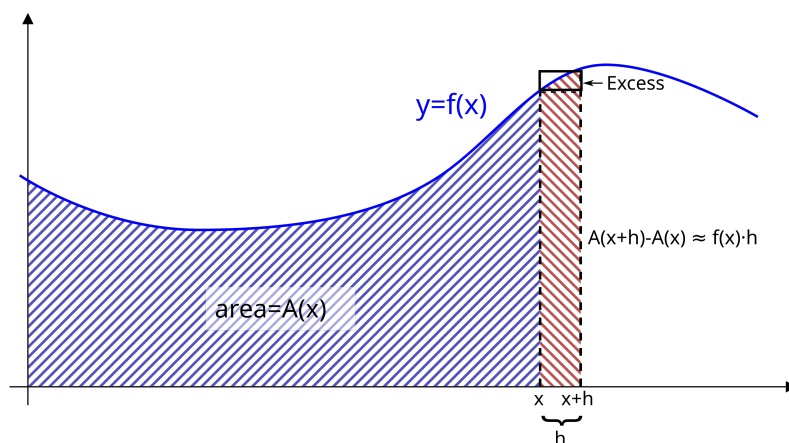


Abbildung 6.2: Die Funktion $F(x)$, in der Figur $A(x)$ genannt, misst die Fläche (englisch: “area”) unter der Funktion $y = f(x)$ bis zur Stelle x . Geht man nun einen Schritt h weiter, so wird die zusätzliche rote Fläche mit der Größe $F(x+h) - F(x) \approx f(x)h$ eingeschlossen. Der Fehler dieser Näherung ist der Überschuss (englisch: “excess”), mit dem die rote Fläche vom Rechteck mit der Fläche $f(x)h$ abweicht. Dividiert man beide Seiten durch h , so erhält man $f(x) \approx \frac{F(x+h)-F(x)}{h}$. Für h gegen 0 wird die Näherung exakt: $f(x) = \lim_{h \rightarrow 0} \frac{F(x+h)-F(x)}{h} = F'(x) = \frac{dF}{dx}$. Dies bedeutet: Die Ableitung der Fläche $F(x)$ unter der Funktion $f(x)$, dh. $F'(x)$, ist gleich der ursprünglichen Funktion $f(x)$. Also ist die Flächenfunktion F die Stammfunktion von f . Differentiation und Integration sind inverse Operationen – dies ist die Essenz des Fundamentalsatzes der Analysis. Bildquelle: Wikipedia.

die Integralfunktion $F : [a, b] \rightarrow \mathbb{R}$ mit

$$F(x) = \int_c^x f(t) dt \quad (6.16)$$

differenzierbar und eine Stammfunktion von f , dh. für alle $x \in [a, b]$ gilt $F'(x) = f(x)$. (Für $x = b$ ist einseitige Differenzierbarkeit gemeint.)

Zweiter Teil bzw zweiter Hauptsatz: Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig mit Stammfunktion $F : [a, b] \rightarrow \mathbb{R}$. Dann gilt die Newton-Leibniz-Formel

$$\int_a^b f(x) dx = F(b) - F(a). \quad (6.17)$$

Für $F(b) - F(a)$ verwendet man oft die kurze Schreibweise $F(x)|_a^b$ oder $[F(x)]_a^b$.

Die Beweisidee und geometrische Intuition diskutieren wir in Figur 6.2. Dort wird gezeigt, dass die Ableitung der Fläche unter einer Funktion gerade die Funktion selbst ist.

Der erste (I) und zweite Teil (II) des Hauptsatzes führen zu folgendem konsistenten Kreisgang:

$$\int_a^b f(x) dx \stackrel{\text{(II)}}{=} F(b) - F(a)$$

$$\begin{aligned}
&\stackrel{\text{(I)}}{=} \int_c^b f(t) \, dt - \int_c^a f(t) \, dt \\
&\stackrel{\text{(II)}}{=} F(b) - F(c) - [F(a) - F(c)] = F(b) - F(a).
\end{aligned} \tag{6.18}$$

Aus der Newton-Leibniz-Formel folgt, dass sich Integrale aufteilen lassen und sich bei Vertauschung der Integrationsgrenzen das Vorzeichen ändert:

$$\int_a^b f(x) \, dx = \int_a^c f(x) \, dx + \int_c^b f(x) \, dx \quad \text{mit } a < c < b, \tag{6.19}$$

$$\int_a^b f(x) \, dx = F(b) - F(a) = - \int_b^a f(x) \, dx. \tag{6.20}$$

Bei der Integration durch Substitution müssen die Integrationsgrenzen entsprechend angepasst werden:

$$\int_a^b f(g(x)) g'(x) \, dx = \int_{g(a)}^{g(b)} f(u) \, du. \tag{6.21}$$

Man kann auch Funktionen integrieren, deren Definitionsbereich unbeschränkt ist. Das Ergebnis bezeichnet man als *uneigentliches Integral*.

Wir betrachten nun ein paar Beispiele, um diese Konzepte zu verinnerlichen:

Beispiel: Berechne das Integral $\int_0^1 x^2 \, dx$ mittels Newton-Leibniz-Formel. *Lösung:*

$$\int_0^1 x^2 \, dx = \left[\frac{x^3}{3} \right]_0^1 = \frac{1}{3} - \frac{0}{3} = \frac{1}{3}. \tag{6.22}$$

Beispiel: Berechne das Integral $\int_{-1}^1 |x| \, dx$ mittels Newton-Leibniz-Formel. *Lösung:*

$$\begin{aligned}
\int_{-1}^1 |x| \, dx &= \int_{-1}^0 |x| \, dx + \int_0^1 |x| \, dx = \int_{-1}^0 (-x) \, dx + \int_0^1 x \, dx \\
&= \left[-\frac{x^2}{2} \right]_{-1}^0 + \left[\frac{x^2}{2} \right]_0^1 = 0 + \frac{1}{2} + \frac{1}{2} - 0 = 1.
\end{aligned} \tag{6.23}$$

Beispiel: Berechne $\int_0^3 x \sin(x^2) \, dx$. *Lösung:* Substituiere $u := x^2$, dh. $du = 2x \, dx$. Damit:

$$\begin{aligned}
\int_0^3 x \sin(x^2) \, dx &= \int_{u(0)=0}^{u(3)=9} \sin(u) \frac{1}{2} \, du \\
&= \frac{1}{2} [-\cos(u)]_0^9 = \frac{1}{2} [-\cos(9) + 1] = 0.955 \dots
\end{aligned} \tag{6.24}$$

Beispiel: Berechne das uneigentliche Integral $\int_1^\infty \frac{1}{x} dx$. *Lösung:*^a

$$\int_1^\infty \frac{1}{x} dx = \lim_{b \rightarrow \infty} \int_1^b \frac{1}{x} dx = \lim_{b \rightarrow \infty} \left[\ln(x) \right]_1^b = \lim_{b \rightarrow \infty} \ln(b) - \ln(1) = \infty. \quad (6.25)$$

^aDies erinnert uns an die Divergenz der harmonischen Reihe $\sum_{k=1}^\infty \frac{1}{k}$.

Beispiel: Berechne das uneigentliche Integral $\int_1^\infty \frac{1}{x^2} dx$. *Lösung:*^a

$$\int_1^\infty \frac{1}{x^2} dx = \lim_{b \rightarrow \infty} \int_1^b \frac{1}{x^2} dx = \lim_{b \rightarrow \infty} \left[-\frac{1}{x} \right]_1^b = \lim_{b \rightarrow \infty} \left[-\frac{1}{b} + \frac{1}{1} \right] = 1. \quad (6.26)$$

^aDies erinnert uns an die Konvergenz der Reihe $\sum_{k=1}^\infty \frac{1}{k^2}$.

6.2 In mehreren Variablen

Wir betrachten nun erneut reellwertige n -dimensionale Funktionen der Form $f(x_1, x_2, \dots, x_n)$. Das n -dimensionale Integral liefert die Stammfunktion

$$F(x_1, x_2, \dots, x_n) = \int \int \cdots \int f(x_1, x_2, \dots, x_n) dx_1 dx_2 \dots dx_n. \quad (6.27)$$

Dies ist so zu verstehen, dass zuerst das innerste Integral ausgeführt wird (und wie beim partiellen Ableiten alle anderen Variablen als konstant betrachtet werden), dann das zweitinnerste, und so weiter, bis zum äußersten.

Satz von Fubini: Für stetige Funktionen f kann bei unbestimmten Integralen oder bestimmten Integralen mit unabhängigen Integrationsgrenzen die Reihenfolge der Integrale vertauscht werden. Wir schreiben dies beispielhaft für eine zweidimensionale Funktion $f(x, y)$ nieder:

$$\int \int f(x, y) dx dy = \int \int f(x, y) dy dx, \quad (6.28)$$

$$\int_c^d \int_a^b f(x, y) dx dy = \int_a^b \int_c^d f(x, y) dy dx. \quad (6.29)$$

Beispiel: Berechne $\int \int xy^2 dx dy$. *Lösung:*

$$\int \int xy^2 dx dy = \int \frac{x^2}{2} y^2 dy = \frac{x^2 y^3}{6}. \quad (6.30)$$

Vertauschen der Integrationsreihenfolge führt zum gleichen Ergebnis: $\int \int xy^2 dy dx = \int x \frac{y^3}{3} dx = \frac{x^2 y^3}{6}$.

Beispiel: Berechne die Fläche eines Rechtecks mit den Seitenlängen a entlang x und b entlang y . *Lösung:* Wir wählen das Rechteck mit seinem linken unteren Eckpunkt in

den Ursprung. Dann ist die Fläche:

$$\int_0^a \int_0^b dy dx = \int_0^a [y]_0^b dx = b \int_0^a dx = b [x]_0^a = ab. \quad (6.31)$$

Vertauschen der Integrationsreihenfolge führt zum gleichen Ergebnis.

Beispiel: Berechne die Fläche A eines Kreises mit Radius R . *Lösung:* Wir betrachten einen Viertelkreis mit Ursprung im Punkt $(0,0)$ im ersten Quadranten, dh. $x \in [0, R]$ und $y \in [0, \sqrt{R^2 - x^2}]$. Das folgende bestimmte Integral – die dafür nötige Stammfunktion findet man mit Computerhilfe oder in Integraltabellen – liefert dann ein Viertel der Kreisfläche, dh. $\frac{A}{4}$:

$$\begin{aligned} \frac{A}{4} &= \int_0^R \int_0^{\sqrt{R^2 - x^2}} 1 dy dx = \int_0^R [y]_0^{\sqrt{R^2 - x^2}} dx = \int_0^R \sqrt{R^2 - x^2} dx \\ &= \left[\frac{x}{2} \sqrt{R^2 - x^2} + \frac{R^2}{2} \arcsin \frac{x}{R} \right]_0^R = \left[0 + \frac{R^2}{2} \arcsin(1) - 0 - \frac{R^2}{2} \arcsin(0) \right] \\ &= \frac{\pi R^2}{4}. \end{aligned} \quad (6.32)$$

In diesem Beispiel darf die Reihenfolge der Integrale nicht vertauscht werden, da die Integrationsgrenzen nicht unabhängig sind. Die innere Grenze hängt von der Integrationsvariable des äußeren Integrals ab.

Koordinatentransformationen

Oft ist es hilfreich, die Symmetrie eines Problems zu berücksichtigen und nicht in kartesischen Koordinaten zu bleiben. Das obige Problem der Kreisfläche wird viel einfacher, wenn wir zu Polarkoordinaten übergehen:

$$x = r \cos \varphi, \quad (6.33)$$

$$y = r \sin \varphi. \quad (6.34)$$

Nun müssen wir noch das Flächenelement $dA = dx dy$ transformieren. Dies funktioniert wie folgt: Wir berechnen zuerst die sogenannte Jacobi-Matrix der ersten Ableitungen:

$$J = \frac{\partial(x,y)}{\partial(r,\varphi)} = \begin{pmatrix} \frac{\partial x}{\partial r} & \frac{\partial x}{\partial \varphi} \\ \frac{\partial y}{\partial r} & \frac{\partial y}{\partial \varphi} \end{pmatrix} = \begin{pmatrix} \cos \varphi & -r \sin \varphi \\ \sin \varphi & r \cos \varphi \end{pmatrix}. \quad (6.35)$$

Und dann davon die Determinante:

$$\det J = r \cos^2 \varphi + r \sin^2 \varphi = r. \quad (6.36)$$

Das Flächenelement transformiert dann wie folgt:

$$dA = dx dy = |\det J| dr d\varphi = r dr d\varphi. \quad (6.37)$$

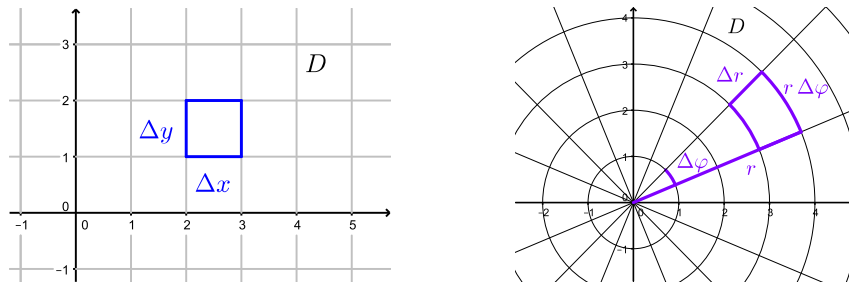


Abbildung 6.3: Links: In kartesischen Koordinaten ist das Flächenelement $dA = dx dy$ (in der Abbildung sind mit Δx und Δy nicht-infinitesimale Größen gezeichnet). Rechts: In Polarkoordinaten ist ein Flächenelement ein Rechteck, wobei die eine Seite durch dr gegeben ist und die andere durch die Bogenlänge $r d\varphi$. Das Flächenelement ist daher $dA = r dr d\varphi$. Bildquelle: *hier*.

Figur 6.3 veranschaulicht diese Transformation.

Beispiel: Berechne die Fläche A eines Kreises K mit Radius R . *Lösung:* In Polarkoordinaten können wir wie folgt integrieren:

$$A = \iint_K dA = \int_0^{2\pi} \int_0^R r dr d\varphi = \int_0^{2\pi} \frac{R^2}{2} d\varphi = \pi R^2. \quad (6.38)$$

Beispiel: Beweise, dass die Gauß-Verteilung (auch genannt: Normalverteilung oder Gaußsche Glockenkurve)

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}} \quad (6.39)$$

eine normierte Wahrscheinlichkeitsverteilung ist, dh. (i) sie ist überall nicht-negativ und (ii) das Integral I über den ganzen Raum gleich 1 ist. Hier sind $\mu \in \mathbb{R}$ der Mittelwert und $\sigma > 0$ die Standardabweichung. *Beweis:* (i) Da die Exponentialfunktion überall positiv ist, gilt $f(x) > 0$ für alle $x \in \mathbb{R}$. (ii) Wir definieren das gesuchte Integral

$$I := \int_{-\infty}^{\infty} f(x) dx = \frac{1}{\sqrt{2\pi}\sigma} \int_{-\infty}^{\infty} e^{-\frac{(x-\mu)^2}{2\sigma^2}} dx. \quad (6.40)$$

Als erstes substituieren wir $u := x - \mu$ mit $du = dx$. Dies führt zu

$$I = \frac{1}{\sqrt{2\pi}\sigma} \int_{-\infty}^{\infty} e^{-\frac{(x-\mu)^2}{2\sigma^2}} dx = \frac{1}{\sqrt{2\pi}\sigma} \int_{u(-\infty)=-\infty}^{u(\infty)=\infty} e^{-\frac{u^2}{2\sigma^2}} du. \quad (6.41)$$

Die (endliche) Verschiebung um μ spielt keine Rolle, da über alle reellen Zahlen integriert wird. Wir können nun $x := u$ substituieren und erhalten

$$I = \frac{1}{\sqrt{2\pi}\sigma} \int_{-\infty}^{\infty} e^{-\frac{x^2}{2\sigma^2}} dx. \quad (6.42)$$

Dieses Integral ist vorerst nicht lösbar. Aber mit folgendem Trick geht es. Wir berechnen zunächst I^2 :

$$\begin{aligned} I^2 &= \left(\frac{1}{\sqrt{2\pi\sigma}} \int_{-\infty}^{\infty} e^{-\frac{x^2}{2\sigma^2}} dx \right)^2 = \frac{1}{\sqrt{2\pi\sigma}} \int_{-\infty}^{\infty} e^{-\frac{x^2}{2\sigma^2}} dx \frac{1}{\sqrt{2\pi\sigma}} \int_{-\infty}^{\infty} e^{-\frac{y^2}{2\sigma^2}} dy \\ &= \frac{1}{2\pi\sigma^2} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-\frac{x^2}{2\sigma^2}} e^{-\frac{y^2}{2\sigma^2}} dx dy = \frac{1}{2\pi\sigma^2} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-\frac{x^2+y^2}{2\sigma^2}} dx dy. \end{aligned} \quad (6.43)$$

Da wir nun im zweidimensionalen Raum sind, können wir auf Polarkoordinaten wechseln, dh. $\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(x,y) dx dy = \int_0^{2\pi} \int_0^{\infty} f(x(r,\varphi), y(r,\varphi)) r dr d\varphi$:

$$\begin{aligned} I^2 &= \frac{1}{2\pi\sigma^2} \int_0^{2\pi} \int_0^{\infty} e^{-\frac{r^2}{2\sigma^2}} r dr d\varphi \\ &= \frac{1}{\sigma^2} \int_0^{\infty} e^{-\frac{r^2}{2\sigma^2}} r dr. \end{aligned} \quad (6.44)$$

Wir substituieren $u := -\frac{r^2}{2\sigma^2}$ mit $du = -\frac{r}{\sigma^2} dr$. Damit folgt:

$$\begin{aligned} I^2 &= \frac{1}{\sigma^2} \int_{u(0)}^{u(\infty)} e^u (-\sigma^2) du = - \int_0^{-\infty} e^u du \\ &= -[e^u]_0^{-\infty} = -[0 - 1] = 1. \end{aligned} \quad (6.45)$$

Nun ziehen wir die Wurzel und erhalten $I = \pm 1$. Die Lösung $I = -1$ ist ausgeschlossen, weil die Funktion $f(x)$ überall positiv ist. Also gilt $I = 1$. $\square$

In drei Raum-Dimensionen sind Zylinder- und Kugelkoordinaten zwei häufig verwendete Koordinatensysteme. Untenstehend notieren wir die entsprechenden Koordinatentransformationen samt Flächenelement dA und Volumenelement dV . Figur 6.4 illustriert die beiden Transformationen.

Kugelkoordinaten ($r \geq 0$, $\theta \in [0, \pi]$, $\varphi \in [0, 2\pi]$):

$$x = r \sin \theta \cos \varphi, \quad (6.46)$$

$$y = r \sin \theta \sin \varphi, \quad (6.47)$$

$$z = r \cos \theta, \quad (6.48)$$

$$dA = r^2 \sin \theta d\theta d\varphi, \quad (6.49)$$

$$dV = r^2 \sin \theta dr d\theta d\varphi. \quad (6.50)$$

Zylinderkoordinaten ($r \geq 0$, $\varphi \in [0, 2\pi]$, $z \in \mathbb{R}$):

$$x = r \cos \varphi, \quad (6.51)$$

$$y = r \sin \varphi, \quad (6.52)$$

$$z = z, \quad (6.53)$$

$$dA = r d\varphi dz, \quad (6.54)$$

$$dV = r dr d\varphi dz. \quad (6.55)$$

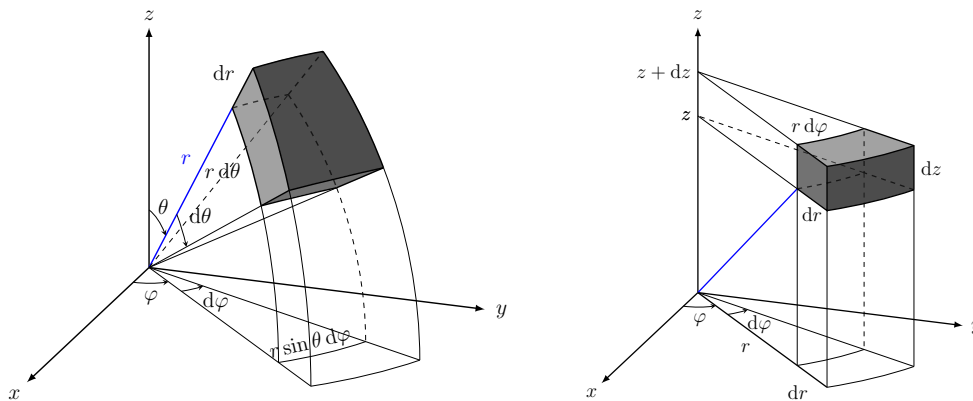


Abbildung 6.4: Links: In Kugelkoordinaten besteht das Volumenelement dV aus einem Quader mit folgenden Kantenlängen: (i) die Änderung entlang r ist das Differential dr , (ii) die Änderung des Polarwinkels $d\theta$ führt zur Bogenlänge $r d\theta$, (iii) die Änderung des Azimuthalwinkels $d\varphi$ führt zur Bogenlänge $r \sin \theta d\varphi$, wobei der Faktor $\sin \theta$ der Tatsache Rechnung trägt, dass man in der Nähe der Pole kleinere Bogenlängen hat als am Äquator. Das Volumenelement ist mithin $dV = r^2 \sin \theta dr d\theta d\varphi$. Das Flächenelement ist der innere Oberflächenteil des Quaders: $dA = r^2 \sin \theta d\theta d\varphi$. Rechts: In Zylinderkoordinaten ist das Flächenelement $dA = r d\varphi dz$ und das Volumenelement ist $dV = r dr d\varphi dz$. Bildquelle: Adaptiert von <https://tikz.net/>

Beispiel: Berechne das Volumen einer Kugel K mit Radius R . *Lösung:* In Kugelkoordinaten erhalten wir

$$V_K = \iiint_K dV = \int_0^{2\pi} \int_0^\pi \int_0^R r^2 \sin \theta dr d\theta d\varphi = \frac{R^3}{3} [-\cos \theta]_0^\pi 2\pi = \frac{4\pi R^3}{3}. \quad (6.56)$$

Beispiel: Berechne die Mantelfläche eines Zylinders Z mit Radius R und Höhe H . *Lösung:* In Zylinderkoordinaten erhalten wir

$$A_Z = \iint_Z dA = \int_0^H \int_0^{2\pi} R d\varphi dz = 2\pi RH. \quad (6.57)$$

6.3 Numerische Integration

Es kommt häufig vor, dass Integrale nicht analytisch lösbar sind. Wir betrachten im folgenden zwei numerische Herangehensweisen, nämlich Quadraturverfahren und Monte-Carlo-Integration.

Quadraturverfahren

Bei der Integration durch Quadratur betrachtet man die zu integrierende Funktion $f(x)$ an $n+1$ Stützstellen $a := x_0 < x_1 < x_2 < \dots < x_n =: b$ und approximiert das Integral durch

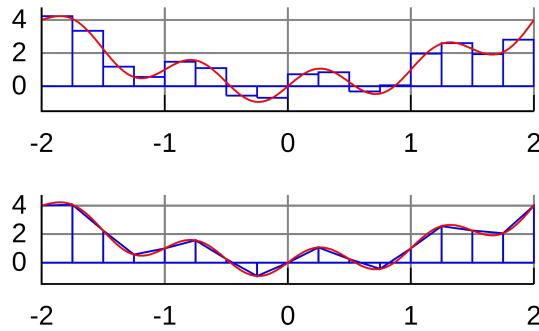


Abbildung 6.5: Quadraturverfahren. Oben: Rechteck-Regel. Unten: Trapez-Regel. Die Funktion $f(x)$ (blau) wird durch Rechtecke bzw. Trapeze (rot) angenähert. Bildquelle: Wikipedia.

eine gewichtete Summe:

$$I(f(x)) := \int_a^b f(x) dx \approx Q(f(x)) := (b-a) \sum_{i=0}^n w_i f(x_i) \quad (6.58)$$

Die Zahl n gibt den Grad der Quadraturformel an, und w_i sind die Gewichte, die im allgemeinen vom Abstand zu den benachbarten Stützstellen abhängen.

Ein einfacher Fall ist die Rechteck-Regel (auch genannt: Mittelpunktsregel). Sie besteht darin, die $n+1$ Stützstellen so zu wählen, dass n äquidistante Intervalle entstehen, dh. $x_i - x_{i-1} = \Delta x = \frac{b-a}{n}$ für $i = 1, 2, \dots, n$. In jedem Intervall $[x_{i-1}, x_i]$ approximiert man die Fläche durch das Rechteck, dessen obere Seite durch den Punkt $(\frac{x_{i-1}+x_i}{2}, f(\frac{x_{i-1}+x_i}{2}))$ geht, also durch den Funktionswert in der Mitte. Diese Fläche ist $\Delta x f(\frac{x_{i-1}+x_i}{2})$. Dann erhält man

$$\begin{aligned} I &= \int_a^b f(x) dx = \int_{x_0}^{x_1} f(x) dx + \int_{x_1}^{x_2} f(x) dx + \dots + \int_{x_{n-1}}^{x_n} f(x) dx \\ &= \sum_{i=1}^n \int_{x_{i-1}}^{x_i} f(x) dx \approx \Delta x \sum_{i=1}^n f\left(\frac{x_{i-1}+x_i}{2}\right) \end{aligned} \quad (6.59)$$

Die Trapez-Regel nähert die Fläche in einem Intervall $[x_{i-1}, x_i]$ durch ein Trapez an, also durch die Fläche unter der Linie von $(x_{i-1}, f(x_{i-1}))$ nach $(x_i, f(x_i))$. Dies ergibt die Approximation

$$I = \int_a^b f(x) dx \approx \Delta x \sum_{i=1}^n \frac{f(x_{i-1}) + f(x_i)}{2}. \quad (6.60)$$

Die Rechteck- und Trapezregel sind in Figur 6.5 veranschaulicht.

Es gibt Methoden, die deutlich besser konvergieren, also mit weniger Stützstellen bessere Genauigkeit erreichen: Die Simpson-Regel legt im $[x_{i-1}, x_i]$ eine Parabel durch die drei Punkte $(x_{i-1}, f(x_{i-1}))$, $(x_{m_i}, f(x_{m_i}))$ und $(x_i, f(x_i))$, wobei $x_{m_i} = \frac{x_{i-1}+x_i}{2}$ der Mittelpunkt zwischen den beiden Stützstellen ist. Weitere Verfahren wie die Gauß-Quadratur verwenden speziell gewählte (nicht äquidistante) Stützstellen und Polynome höherer Ordnung.

Quadraturverfahren können (mit entsprechendem Aufwand) auf mehrere Dimensionen

verallgemeinert werden.

Monte-Carlo-Integration

Das nach dem Casino in Monte Carlo benannte Verfahren beruht darauf, n zufällige Punkte $x_1, x_2, \dots, x_n$ gleichverteilt im Integrationsintervall zu ziehen und dann das Integral durch das arithmetische Mittel der Funktionswerte an diesen Stellen anzunähern:

$$I(f(x)) := \int_a^b f(x) \, dx \approx \frac{b-a}{n} \sum_{i=1}^n f(x_i). \quad (6.61)$$

Dieses Verfahren ist besonders gut auf höhere Dimensionen erweiterbar und funktioniert dort oft besser als klassische Integrationsalgorithmen.

7. Differentialgleichungen

Differentialgleichungen sind Gleichungen, in denen eine Funktion in ihrer Ableitung bzw. in mehreren ihrer höheren Ableitungen vorkommt. Sie sind besonders relevant in Naturwissenschaft und Technik und beschreiben eine Vielzahl von natürlichen oder designten Prozessen. Konkret nehmen viele Naturgesetze die Gestalt von Differentialgleichungen an, wie etwa das zweite Newtonsche Gesetz (Kraftgesetz) der klassischen Mechanik oder die Schrödinger-Gleichung der Quantenmechanik.

Für manche Differentialgleichungen gibt es Lösungsmethoden wie zB. Trennung der Variablen (siehe unterhalb) oder mittels Laplace- oder Fourier-Transformation. In vielen Fällen können Differentialgleichungen aber nicht explizit analytisch gelöst werden und man muss auf numerische Methoden zurückgreifen.

7.1 Gewöhnliche Differentialgleichungen

Eine Differentialgleichung, in der eine Funktion $y(x)$ und deren Ableitungen nach nur einer Variablen x vorkommen, nennt man *gewöhnliche Differentialgleichung*. Gewöhnliche Differentialgleichungen haben die Form

$$F(x, y(x), y'(x), y''(x), \dots, y^{(n)}(x)) = 0. \quad (7.1)$$

Der höchste vorkommende Grad n der Ableitung gibt die *Ordnung* der Differentialgleichung an.

Beispiel: Die Gleichung

$$y'''(x) + 2y'(x) - y^2(x) + \sin(x) + c = 0 \quad (7.2)$$

ist eine gewöhnliche Differentialgleichung 3. Ordnung. Sie ist darüber hinaus nicht-linear, da das Quadrat von y vorkommt. Der Term $\sin(x)$ tut bezüglich Linearität nichts zur Sache, da sich die Linearität nur auf die gesuchte Funktion y bezieht, nicht aber auf

die Variable x .

Anfangswertproblem

Anfangswertprobleme haben ihre Motivation in der Physik, wobei die gesuchte Funktion y von der Zeit t abhängig ist. Wir schreiben nun also $y(t)$, und nicht mehr $y(x)$. Ein Anfangswertproblem erster Ordnung ist eine Differentialgleichung erster Ordnung mit einer zusätzlichen Anfangsbedingung zum Zeitpunkt $t = 0$:

$$y'(t) = f(t, y(t)), \quad y(0) = y_0. \quad (7.3)$$

Wenn f stetig ist, erhalten wir die formale Lösung mithilfe des ersten Hauptsatzes der Differential- und Integralrechnung wie folgt:

$$y(t) = y_0 + \int_0^t f(t', y(t')) dt'. \quad (7.4)$$

Im allgemeinen Fall ist dieses Integral aber nicht analytisch ausführbar, weil die gesuchte Funktion $y(t)$ im Integral vorkommt.

Ein Anfangswertproblem k -ter Ordnung kann in ein System von Anfangswertproblemen erster Ordnung überführt werden.

Beispiel: Der radioaktive Zerfall: Zu jedem Zeitpunkt t gibt es von einem Material $N(t)$ Teilchen. Die Änderung der Teilchenzahl, dh. die erste Ableitung von N nach t , zu einem Zeitpunkt t ist abhängig von der Zahl der zu ebendiesem Zeitpunkt noch vorhandenen Teilchen. (Je mehr Teilchen es gibt, umso mehr zerfallen.) Zum Start-Zeitpunkt $t = 0$ sind genau N_0 Teilchen vorhanden. Der Prozess ist ein Anfangswertproblem erster Ordnung:

$$N'(t) \equiv \frac{dN(t)}{dt} = -\lambda N(t), \quad N(0) = N_0. \quad (7.5)$$

Hier ist λ die Zerfallskonstante, die vom Material abhängt. Das Minuszeichen stellt sicher, dass die Teilchenzahl abnimmt und nicht wächst. Die formale Lösung

$$N(t) = N_0 + \int_0^t (-\lambda) N(t') dt' \quad (7.6)$$

hilft uns nicht weiter, da die gesuchte Funktion im Integral auftaucht. Die Lösung von (7.5) erfolgt durch *Trennung der Variablen* – dh. in diesem Fall, dass alle N auf eine Seite und alle t auf die andere gebracht werden – und anschließende Integration beider Seiten:

$$\frac{dN}{N} = -\lambda dt, \quad (7.7)$$

$$\int \frac{dN}{N} = - \int \lambda dt, \quad (7.8)$$

$$\ln N = -\lambda t + \bar{c}, \quad (7.9)$$

$$N(t) = e^{-\lambda t + \bar{c}} = c e^{-\lambda t}. \quad (7.10)$$

Die Konstante c kann man über die Anfangsbedingung finden: $N(0) = c = N_0$. Das führt zur vollständigen Lösung des radioaktiven Zerfalls:

$$N(t) = N_0 e^{-\lambda t}. \quad (7.11)$$

Zum Zeitpunkt $t = \frac{1}{\lambda}$ sind noch $N_0 \frac{1}{e}$, dh. $\approx 37\%$ der anfänglichen Teilchen vorhanden.

Beispiel: Der freie Fall: Die auf einen Körper mit Masse m einwirkende Schwerkraft $F = mg$ mit $g \approx 9.81 \text{ ms}^{-2}$ der Schwerkbeschleunigung an der Erdoberfläche ruft eine Beschleunigung a hervor, wobei $F = ma$ gilt (zweites Newtonsches Gesetz). Betrachten wir den eindimensionalen Fall, in dem der Ort eines Teilchens zu einem Zeitpunkt t durch $x(t)$ gegeben ist, wobei die x -Achse zum Erdmittelpunkt zeigt. Die zeitliche Ableitung des Orts ist die Geschwindigkeit $v(t) = \frac{dx(t)}{dt}$. Die Ableitung der Geschwindigkeit ist die Beschleunigung $a = \frac{dv(t)}{dt} = \frac{d^2x(t)}{dt^2}$. Der freie Fall hat somit die Form einer gewöhnlichen Differentialgleichung zweiter Ordnung

$$m \frac{d^2x(t)}{dt^2} = mg. \quad (7.12)$$

Dies können wir in ein System aus zwei Gleichungen erster Ordnung umschreiben:

$$\frac{dv(t)}{dt} = g, \quad (7.13)$$

$$\frac{dx(t)}{dt} = v(t). \quad (7.14)$$

Die erste Gleichung (7.13) kann direkt durch Integration (erster Hauptsatz) gelöst werden:

$$v(t) = v_0 + \int_0^t \frac{dv(t')}{dt'} dt' = v_0 + \int_0^t g dt' = v_0 + gt. \quad (7.15)$$

Nun lösen wir die zweite Gleichung (7.14) unter Verwendung der Lösung (7.15). Aus didaktischen Gründen verwenden wir nicht den ersten Hauptsatz sondern integrieren hier erst nach Trennung der Variablen:

$$dx = v(t) dt, \quad (7.16)$$

$$dx = (v_0 + gt) dt, \quad (7.17)$$

$$\int dx = \int (v_0 + gt) dt, \quad (7.18)$$

$$x(t) = x_0 + v_0 t + \frac{1}{2} g t^2. \quad (7.19)$$

Die beiden Integrationskonstanten v_0 und x_0 sind die Geschwindigkeit bzw. der Ort zum Zeitpunkt $t = 0$, die man aus den Anfangsbedingungen erhält: $\frac{dx}{dt}(0) = v(0) = v_0$ und $x(0) = x_0$.

Beispiel: Der harmonische Oszillator bzw. das Federpendel: Eine Masse m ist an einer Feder befestigt, die eine Kraft entgegen der Auslenkung x um die Null-Lage erzeugt. Die Federkonstante k gibt die Steifigkeit der Feder an. Die gewöhnliche Differentialgleichung zweiter Ordnung ist

$$m \frac{d^2 x(t)}{dt^2} = -kx(t). \quad (7.20)$$

Lösen durch Trennung der Variablen und Integration ist hier nicht möglich. Diese Gleichung wird über einen sogenannten *Ansatz* gelöst, dh. man rät die Gestalt der Lösung mit offenen Parametern. Wir wissen, dass eine Feder zu einer Schwingung führt. Außerdem wissen wir, dass die zweifache Ableitung der Sinus-Funktion wieder die (negative) Sinus-Funktion ergibt. Also vermuten wir die Lösung

$$x(t) = A \sin(\omega t + \varphi), \quad (7.21)$$

wobei A die Amplitude, ω die Frequenz und φ eine Phasenverschiebung ist. Wir erhalten die Ableitungen

$$\frac{dx(t)}{dt} = A \omega \cos(\omega t + \varphi), \quad (7.22)$$

$$\frac{d^2 x(t)}{dt^2} = -A \omega^2 \sin(\omega t + \varphi). \quad (7.23)$$

Einsetzen in (7.20) führt zu

$$-mA \omega^2 \sin(\omega t + \varphi) = -kA \sin(\omega t + \varphi). \quad (7.24)$$

Diese Gleichung ist erfüllt für $\omega = \sqrt{\frac{k}{m}}$. Also war der Ansatz richtig. Die Konstanten A und φ erhält man aus der Anfangsposition $x(0)$ und der Anfangsgeschwindigkeit $\frac{dx}{dt}(0)$.

7.2 Partielle Differentialgleichungen

Partielle Differentialgleichungen enthalten, wie der Name schon sagt, partielle Ableitungen. Für lineare Gleichungen sind oftmals Lösungen bekannt, nichtlineare müssen zumeist numerisch gelöst werden.

Beispiel: Die Wellengleichung: In einer räumlichen Dimension wird das Verhalten einer Welle mit der Amplitude $u(x, t)$, wobei x die Raumkoordinate und t die Zeit ist, durch

eine partielle Differentialgleichung zweiter Ordnung beschrieben:

$$\frac{\partial^2 u(x,t)}{\partial t^2} = c^2 \frac{\partial^2 u(x,t)}{\partial x^2}. \quad (7.25)$$

Hier ist c die Ausbreitungsgeschwindigkeit der Welle. Die allgemeine Lösung dieser Wellengleichung ist

$$u(x,t) = f(x+ct) + g(x-ct), \quad (7.26)$$

wobei f und g beliebige zweimal ableitbare Funktionen sind. Hier ist $f(x+ct)$ eine nach links laufende und $g(x-ct)$ eine nach rechts laufende Welle. Wir überprüfen/beweisen die Lösung:

$$\frac{\partial u(x,t)}{\partial t} = c f'(x+ct) - c g'(x-ct), \quad (7.27)$$

$$\frac{\partial^2 u(x,t)}{\partial t^2} = c^2 f''(x+ct) + c^2 g''(x-ct), \quad (7.28)$$

$$\frac{\partial u(x,t)}{\partial x} = f'(x+ct) + g'(x-ct), \quad (7.29)$$

$$\frac{\partial^2 u(x,t)}{\partial x^2} = f''(x+ct) + g''(x-ct). \quad (7.30)$$

Beispiel: Die Wärmeleitungsgleichung: Die Entwicklung der Temperatur $u(x,t)$ an der Stelle x und zum Zeitpunkt t in einem homogenen Medium mit der Temperaturleitfähigkeit a wird durch eine partielle Differentialgleichung zweiter Ordnung beschrieben:

$$\frac{\partial u(x,t)}{\partial t} = a \frac{\partial^2 u(x,t)}{\partial x^2}. \quad (7.31)$$

Wir benötigen also eine Funktion, deren einfache Ableitung nach t bis auf eine Konstante gleich ist der zweifachen Ableitung nach x . Wir machen daher den Ansatz

$$u(x,t) = \frac{1}{\sqrt{t}} e^{-c \frac{x^2}{t}}. \quad (7.32)$$

Und bilden die Ableitungen:

$$\frac{\partial u(x,t)}{\partial t} = -\frac{\exp(-c \frac{x^2}{t}) (t - 2cx^2)}{2t^{\frac{5}{2}}}, \quad (7.33)$$

$$\frac{\partial^2 u(x,t)}{\partial x^2} = -\frac{2c \exp(-c \frac{x^2}{t}) (t - 2cx^2)}{t^{\frac{5}{2}}}. \quad (7.34)$$

Mit $c = \frac{1}{4a}$ ist die Wärmeleitungsgleichung erfüllt.

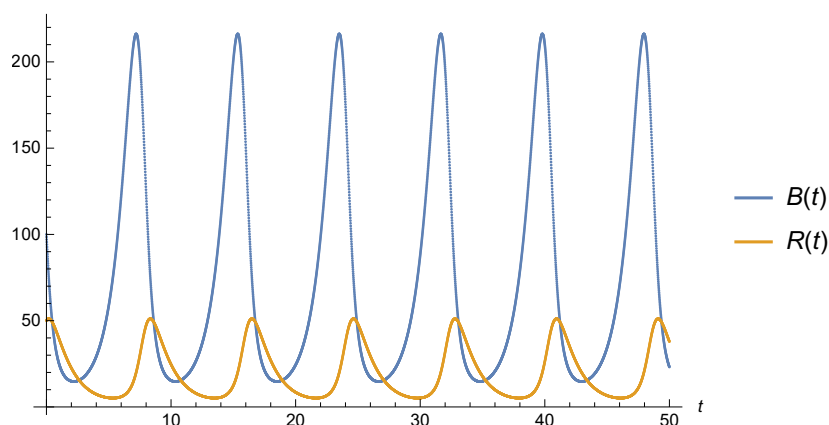


Abbildung 7.1: Numerische Lösung des Räuber-Beute-Modells mithilfe des verbesserten Euler-Verfahrens mit Schrittweite $h = 0.01$. Anfangswerte: $B(0) = 100$, $R(0) = 50$. Modell-Parameter: $r_B = 1$, $f_B = 0.05$, $s_R = 0.75$, $r_R = 0.01$.

7.3 Systeme von Differentialgleichungen

Wenn gleichzeitig mehrere Differentialgleichungen gelöst werden müssen, spricht man von einem *Differentialgleichungssystem* bzw. von *gekoppelten Differentialgleichungen*.

Beispiel: Das Räuber-Beute-Modell (Lotka-Volterra-Gleichungen): Es gibt eine Population $B(t)$ von Beute-Lebewesen, wobei t wieder die Zeit ist. Ihre Reproduktionsrate ist r_B . Die Fressrate pro Räuber und pro Beute-Lebewesen ist f_B . Gleichzeitig gibt es eine Population $R(t)$ von Räubern. Diese haben eine Reproduktionsrate r_R pro Beute-Lebewesen und Räuber. Und eine Sterberate s_R . Die beiden Populationen entwickeln sich entsprechend dem folgenden System aus Differentialgleichungen:

$$\frac{dB}{dt} = r_B B - f_B R B, \quad (7.35)$$

$$\frac{dR}{dt} = -s_R R + r_R B R. \quad (7.36)$$

Je mehr Beute-Lebewesen es gibt, umso besser können sich die Räuber vermehren. Aber je mehr Räuber es gibt, umso mehr Beute-Lebewesen werden gefressen. Dies führt zu Oszillationen, wobei die Räuber zeitlich verzögert (dh. phasenverschoben) sind. Dieses Modell bildet eine wesentliche Grundlage der theoretischen Biologie, insbesondere der Populationsdynamik.

Das Räuber-Beute-Modell kann im allgemeinen nur numerisch gelöst werden. Figur 7.1 zeigt die Lösung für eine spezifische Wahl der Parameter.

7.4 Numerische Verfahren

Für Differentialgleichungen gibt es viele numerische Lösungsmethoden. Wir konzentrieren uns hier auf das sehr intuitive sogenannte (explizite) *Euler-Verfahren*, gefunden von Leonhard Euler im Jahr 1768.

Wir betrachten das Anfangswertproblem einer gewöhnlichen Differentialgleichung erster Ordnung:

$$\frac{dy}{dt} = f(t, y), \quad y(t_0) = y_0. \quad (7.37)$$

Nun wählen wir eine Schrittweite $h > 0$ und damit die diskreten Zeitpunkte

$$t_k = t_0 + kh, \quad k = 0, 1, 2, \dots \quad (7.38)$$

Das *Euler-Verfahren* liefert den (annähernden) Funktionswert y_{k+1} zum Zeitpunkt t_{k+1} aus dem Wert y_k zur Zeit t_k durch die folgende rekursive Beziehung:

$$y_{k+1} = y_k + h f(t_k, y_k). \quad (7.39)$$

Die Werte y_k nähern die tatsächlichen Werte $y(t_k)$ umso besser an, je kleiner die Schrittweite h ist. Aber je kleiner h ist, umso länger dauert es natürlich, bis zu einer bestimmten Zeit zu rechnen. Hergeleitet wird das Euler-Verfahren via Quadratur:

$$\begin{aligned} y(t_{k+1}) &= y(t_k) + \int_{t_k}^{t_{k+1}} f(t, y(t)) dt \\ &\approx y_k + \int_{t_k}^{t_{k+1}} f(t_k, y_k) dt \\ &= y_k + h f(t_k, y_k) =: y_{k+1}. \end{aligned} \quad (7.40)$$

Dh. das Integral von t_k nach t_{k+1} wird ersetzt durch ein Rechteck mit der Breite h und der Höhe $f(t_k, y_k)$, die dem Funktionswert am Intervall-Anfang bei t_k, y_k entspricht (linksseitige Box-Regel). Der Approximationsfehler in jedem Schritt heißt Diskretisierungsfehler. Figur 7.2 veranschaulicht das Euler-Verfahren.

Den Diskretisierungsfehler des Euler-Verfahrens kann man deutlich verkleinern, indem man von der linksseitigen Box-Regel zur Rechteckregel (Mittelpunktsregel) übergeht, also in der Quadratur den Funktionswert von f nicht am Intervall-Anfang sondern in der Intervall-Mitte verwendet. Dies ergibt

Das *verbesserte Euler-Verfahren*: Für jeden Update-Schritt verwendet man ein Rechteck mit der Breite h , aber mit der Höhe des Werts im Mittelpunkt des Intervalls. Um den Mittelpunkt-Wert $f(t_k + \frac{h}{2}, y_{k+\frac{1}{2}})$ zu berechnen, schätzt man den Funktionswert $y_{k+\frac{1}{2}}$ mithilfe des (normalen) Euler-Verfahrens mit Zeitschritt $\frac{h}{2}$:

$$y_{k+1} = y_k + h f(t_k + \frac{h}{2}, y_{k+\frac{1}{2}}) \quad \text{mit} \quad y_{k+\frac{1}{2}} = y_k + \frac{h}{2} f(t_k, y_k). \quad (7.41)$$

Falls wir zwei gekoppelte Differentialgleichungen haben der Form

$$\frac{dy}{dt} = f(t, y, z), \quad y(t_0) = y_0, \quad (7.42)$$

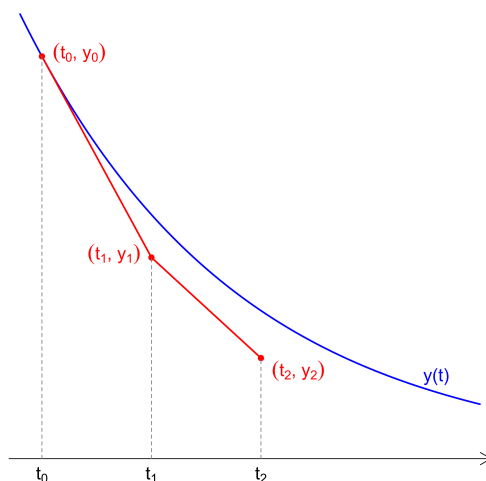


Abbildung 7.2: Die ersten beiden Schritte des Euler-Verfahrens für die Differentialgleichung $\frac{dy}{dt} = f(t, y)$: Im Startpunkt (t_0, y_0) wird die Ableitung von y ermittelt und dann einen Schritt h in diese Richtung gegangen zum Punkt (t_1, y_1) . Dies wird für alle weiteren Schritte entsprechend wiederholt. Die tatsächliche Funktion $y(t)$ (blau) wird durch das Euler-Verfahren (rot) approximiert. Bildquelle: Wikipedia.

$$\frac{dz}{dt} = g(t, y, z), \quad z(t_0) = z_0, \quad (7.43)$$

dann liefert das Euler-Verfahren die Lösungen

$$y_{k+1} = y_k + h f(t_k, y_k, z_k), \quad (7.44)$$

$$z_{k+1} = z_k + h g(t_k, y_k, z_k). \quad (7.45)$$

Und das verbesserte Euler-Verfahren führt zu

$$y_{k+1} = y_k + h f\left(t_k + \frac{h}{2}, y_k + \frac{h}{2} f(t_k, y_k, z_k), z_k + \frac{h}{2} g(t_k, y_k, z_k)\right), \quad (7.46)$$

$$z_{k+1} = z_k + h g\left(t_k + \frac{h}{2}, y_k + \frac{h}{2} f(t_k, y_k, z_k), z_k + \frac{h}{2} g(t_k, y_k, z_k)\right). \quad (7.47)$$

Beispiel: Lösen Sie das Räuber-Beute-Modell mittels verbessertem Euler-Verfahren und reproduzieren Sie damit Figur 7.1. *Lösung:* In den Übungen.

Jenseits des verbesserten Euler-Verfahrens gibt es noch deutlich kompliziertere Methoden, die noch kleinere Diskretisierungsfehler liefern, indem bei der Berechnung eines Schritts mehrere zuvor berechnete Werte verwendet werden (Mehrschritt-Verfahren) oder die Funktion an mehreren Stellen des Schrittingintervalls einbezogen wird (Runge-Kutta-Verfahren).

A. Formelsammlung

Die nachstehende Formelsammlung wird im Anhang der Klausur-Angaben abgedruckt sein. Bei der Klausur selbst sind nur diese weiteren Materialien erlaubt: Papier, Stifte, Radierer, Lineal. Insbesondere dürfen keine Skripten oder Bücher und keine technischen Hilfsmittel (Laptop, Handy, Taschenrechner und ähnliches) verwendet werden.

Reelle und komplexe Zahlen

- Körperaxiome: Kommutativgesetze: $a + b = b + a$ und $ab = ba$, Assoziativgesetze: $a + (b + c) = (a + b) + c$ und $a(bc) = (ab)c$, Distributivgesetz: $a(b + c) = ab + ac$, Existenz neutraler Elemente: $a + 0 = a$ und $a \cdot 1 = a$, Existenz inverser Elemente: $a + (-a) = 0$ und (für $a \neq 0$;) $a \cdot a^{-1} = 1$.
- Umgekehrte Dreiecksungleichung und Dreiecksungleichung:
$$|x| - |y| \leq |x \pm y| \leq |x| + |y|$$
- Eulersche Formel: $e^{i\varphi} = \cos \varphi + i \sin \varphi$
- Wurzeln: Die Lösungen von $z^n = |c|e^{i\phi}$ sind $z_k = \sqrt[n]{|c|} \exp\left(\frac{i\phi}{n} + k\frac{2\pi i}{n}\right)$, $k = 0, 1, 2, \dots, n-1$.

Folgen und Reihen

- Konvergenz: Die Folge $(a_n)_{n \in \mathbb{N}}$ konvergiert gegen a für $n \rightarrow \infty$ falls
$$\forall \varepsilon > 0 \exists N_\varepsilon \in \mathbb{N} \forall n \geq N_\varepsilon : |a_n - a| < \varepsilon,$$
- Monotonie-Kriterium: Eine monoton wachsende Folge $(a_n)_{n \in \mathbb{N}}$ konvergiert genau dann, wenn sie nach oben beschränkt ist.
- Cauchy-Kriterium: Eine Folge $(a_n)_{n \in \mathbb{N}}$ heißt Cauchy-Folge, wenn gilt
$$\forall \varepsilon > 0 \exists N_\varepsilon \in \mathbb{N} \forall n, m \geq N_\varepsilon : |a_n - a_m| < \varepsilon$$

Eine Folge konvergiert genau dann, wenn sie eine Cauchy-Folge ist. Eine Reihe konvergiert genau dann, wenn ihre Partialsummen eine Cauchy-Folge bilden.
- Nullfolgenkriterium: Wenn a_k keine Nullfolge ist, dann divergiert die Reihe $\sum_{k=1}^n a_k$.
- Leibniz-Kriterium: Wenn a_k eine Nullfolge ist, dann konvergiert die Reihe $\sum_{k=1}^n (-1)^k a_k$.
- Vergleichskriterium: Gegeben $\forall k \in \mathbb{N} : 0 \leq a_k \leq b_k$. Majorantenkriterium: Wenn $\sum_{k=1}^\infty b_k$ konvergiert, dann konvergiert $\sum_{k=1}^\infty a_k$. Minorantenkriterium: Wenn $\sum_{k=1}^\infty a_k$ divergiert, dann divergiert $\sum_{k=1}^\infty b_k$.
- Quotientenkriterium: $\sum_{k=1}^n a_k$ konvergiert (bzw. divergiert), wenn $\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| < 1$ (bzw. > 1).
- Wurzelkriterium: $\sum_{k=1}^n a_k$ konvergiert (bzw. divergiert), wenn $\limsup_{k \rightarrow \infty} \sqrt[k]{|a_k|} < 1$ (bzw. > 1).
- Potenzreihe: $f(x) = \sum_{k=0}^\infty a_k x^k$, Konvergenzradius: $r = \frac{1}{L}$ mit $L = \limsup_{k \rightarrow \infty} \sqrt[k]{|a_k|}$ oder $\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right|$

Stetigkeit

- Stetigkeit: f ist stetig in x_0 , falls für jede Nullfolge $(h_n)_{n \in \mathbb{N}}$ gilt: $\lim_{n \rightarrow \infty} f(x_0 + h_n) = f(x_0)$.
- Epsilon-Delta-Kriterium: f ist stetig in x_0 , falls $\forall \varepsilon > 0 \exists \delta > 0 : |x - x_0| < \delta \Rightarrow |f(x) - f(x_0)| < \varepsilon$.

Differentialrechnung

- Differenzierbarkeit: f ist differenzierbar in x_0 , falls für jede Nullfolge $(h_n)_{n \in \mathbb{N}}$ der Grenzwert $\lim_{n \rightarrow \infty} \frac{f(x_0 + h_n) - f(x_0)}{h_n}$ existiert.
- Produktregel: $(fg)'(x) = f'(x)g(x) + f(x)g'(x)$
- Quotientenregel: $\left(\frac{f}{g}\right)'(x) = \frac{f'(x)g(x) - f(x)g'(x)}{g^2(x)}$
- Kettenregel: $[g(f(x))]' = g'(f(x))f'(x)$
- Umkehrregel: $(f^{-1})'(x) = \frac{1}{f'(f^{-1}(x))}$
- Taylor-Reihe:

$$f(x) = \sum_{n=0}^{\infty} \frac{f^{(n)}(x_0)}{n!} (x - x_0)^n$$

- Gradient (in zwei Dimensionen):

$$\nabla f(x_1, x_2) = \begin{pmatrix} \frac{\partial f(x_1, x_2)}{\partial x_1} \\ \frac{\partial f(x_1, x_2)}{\partial x_2} \end{pmatrix}$$

- Hesse-Matrix (in zwei Dimensionen):

$$H_f(x_1, x_2) = \begin{pmatrix} \frac{\partial^2 f(x_1, x_2)}{\partial x_1^2} & \frac{\partial^2 f(x_1, x_2)}{\partial x_1 \partial x_2} \\ \frac{\partial^2 f(x_1, x_2)}{\partial x_2 \partial x_1} & \frac{\partial^2 f(x_1, x_2)}{\partial x_2^2} \end{pmatrix}$$

Integralrechnung

- Partielle Integration: $\int f'(x)g(x) dx = f(x)g(x) - \int f(x)g'(x) dx$
- Integration durch Substitution: $\int f(g(x))g'(x) dx = \int f(u) du|_{u=g(x)}$
- Hauptsatz der Differential- und Integralrechnung:

$$1. F(x) = \int_c^x f(t) dt$$

$$2. \int_a^b f(x) dx = F(b) - F(a)$$

- Polarkoordinaten ($r \geq 0, \varphi \in [0, 2\pi[$):

$$x = r \cos \varphi, y = r \sin \varphi$$

$$dA = r dr d\varphi$$

- Kugelkoordinaten ($r \geq 0, \theta \in [0, \pi[, \varphi \in [0, 2\pi[$):

$$x = r \sin \theta \cos \varphi, y = r \sin \theta \sin \varphi, z = r \cos \theta$$

$$dA = r^2 \sin \theta d\theta d\varphi, dV = r^2 \sin \theta dr d\theta d\varphi$$

- Zylinderkoordinaten ($r \geq 0, \varphi \in [0, 2\pi[, z \in \mathbb{R}$):

$$x = r \cos \varphi, y = r \sin \varphi, z = z$$

$$dA = r d\varphi dz, dV = r dr d\varphi dz$$