

Skript zum
Modul Mathematik

geeignet für

Bachelor of Science Biologie und Chemie
und technische Studiengänge

Universitäten Heidelberg und Freiburg

Version 3.1

Moritz Diehl, Torsten Fischer, Dirk Lebiedz

unter Mithilfe von

Kai Antweiler, Melanie Franzem, Lorenz Steinbock und Kristian Wadel

21. Februar 2010

Inhaltsverzeichnis

Einführung	9
1 Einführung in die mathematische Logik	17
1.1 Aussagen und logische Verknüpfungen	17
1.2 Aussageformen und Quantoren	20
1.3 Wahre Aussagen in der Mathematik	22
1.4 Vollständige Induktion	24
1.4.1 Induktion und Deduktion	24
1.4.2 Technik der vollständigen Induktion	25
1.5 Binomialkoeffizient und binomischer Lehrsatz	27
2 Lineare Algebra I	31
2.1 Mengen und Abbildungen	32
2.1.1 Mengen	32
2.1.2 Das kartesische Produkt	32
2.1.3 Abbildungen	33
2.2 Reelle Vektorräume	35
2.2.1 Der $\mathbb{R}^n$ als reeller Vektorraum	35
2.2.2 Allgemeine Vektorräume	37
2.2.3 Untervektorräume	37
2.3 *Gruppen, Körper, Vektorräume	38
2.3.1 Gruppen	38
2.3.2 Körper	40
2.3.3 Allgemeine Vektorräume	40
2.4 Skalarprodukt, euklidische Norm und Vektorprodukt	42
2.4.1 Norm und Distanz	42
2.4.2 Eigenschaften des Skalarproduktes	43
2.4.3 Das Vektorprodukt	45
2.5 Lineare Unabhängigkeit, Basis und Dimension	45
2.5.1 Basis-Isomorphismen	48
2.6 Lineare Abbildungen	49
2.6.1 Bild, Rang und Kern	51
2.7 Matrizen	52

2.7.1	Rechenregeln für Matrizen	54
2.7.2	Von der Matrix zur linearen Abbildung	55
2.7.3	Inversion von Matrizen	56
2.8	Lineare Gleichungssysteme	58
2.8.1	Homogene lineare Gleichungssysteme	58
2.8.2	Inhomogene lineare Gleichungssysteme	63
2.8.3	Praktisches Lösungsverfahren	65
3	Lineare Algebra II	69
3.1	Determinanten	69
3.1.1	Determinante einer (2×2) -Matrix	69
3.1.2	*Permutationen	72
3.1.3	Eigenschaften der Determinante	74
3.1.4	Praktische Berechnung von Determinanten	77
3.2	Eigenwerte und Eigenvektoren	78
3.2.1	Definition von Eigenwerten und Eigenvektoren	84
3.3	Basen und Koordinatensysteme	86
3.3.1	Koordinatentransformation für lineare Abbildungen	92
3.3.2	Ähnlichkeit von Matrizen	94
3.3.3	Diagonalisierbarkeit	95
4	Analysis über $\mathbb{R}$	99
4.1	Folgen und Konvergenz	99
4.2	Teilfolgen	105
4.2.1	*Der Satz von Bolzano-Weierstraß	106
4.2.2	*Limes inferior und Limes superior	108
4.3	Reihen	109
4.3.1	Konvergenzkriterien für Reihen	110
4.3.2	*Alternierende Reihen	111
4.3.3	*Absolute Konvergenz	112
4.4	Exponentialfunktion und Logarithmus	114
4.4.1	Eigenschaften der Exponentialfunktion	114
4.4.2	Der natürliche Logarithmus	117
4.4.3	Potenzen und Logarithmen zu einer positiven Basis	118
4.5	Stetigkeit	118
4.6	Differenzierbarkeit	125
4.7	Der Mittelwertsatz	135
4.8	Taylorentwicklung	140
4.9	Maxima und Minima	143
4.9.1	*Eine Optimierungsaufgabe	145
4.10	Integration	146
4.10.1	*Definition des Riemann-Integrals	148
4.10.2	Einige Sätze zum Integral	151

4.10.3	Rechenregeln zur Integration	153
4.10.4	Uneigentliche Integrale	159
4.10.5	*Maß- und Integrationstheorie (Lebesgue-Integral)	163
4.11	Numerische Differentiation und Integration	165
4.11.1	Numerische Berechnung von Ableitungen mit Hilfe von „Finiten Differenzen“-Approximationen (sog. Diskretisierungen)	165
4.11.2	Numerische Integration	167
5	Komplexe Zahlen	169
5.1	Definition der Menge der komplexen Zahlen	169
5.2	Rechenregeln	171
5.3	Überblick über Zahlbereiche und deren Strukturen	174
6	Skalarprodukte und Orthogonalität	177
6.1	Standard-Skalarprodukt in $\mathbb{R}^n$	177
6.2	Orthogonale Projektion auf eine Gerade	178
6.3	Orthogonale Projektion auf einen Unterraum	182
6.4	Skalarprodukte auf reellen Vektorräumen	183
6.5	Fourier-Entwicklung	185
6.6	*Orthonormalbasen und Selbstadjungierte Operatoren	192
6.6.1	Orthonormalbasen und Orthogonale Matrizen	192
6.6.2	Selbstadjungierte Operatoren und Symmetrische Matrizen	195
6.6.3	*Verallgemeinerung auf komplexe Matrizen	197
6.6.4	Der Laplace-Operator	198
7	Analysis im $\mathbb{R}^n$	203
7.1	Kurven	204
7.1.1	Wie berechnet man die Kurvenlänge?	205
7.2	Ableitungen im $\mathbb{R}^n$	208
7.2.1	Veranschaulichung von Funktionen mehrerer Argumente	208
7.2.2	*Offene Mengen	210
7.3	Allgemeines Konzept der Ableitung	211
7.4	Konvergenz und Stetigkeit	213
7.4.1	Partielle Ableitungen	214
7.4.2	Totale Ableitung	215
7.4.3	Vollständiges Differential	219
7.4.4	Partielle Ableitungen höherer Ordnung	221
7.5	Taylor-Entwicklung und lokale Extrema	222
7.6	Funktionen vom $\mathbb{R}^n$ in den $\mathbb{R}^m$	226
7.6.1	Allgemeiner Ableitungsbegriff im $\mathbb{R}^n$	226
7.7	Integration im $\mathbb{R}^n$	229
7.7.1	Sukzessive Integration	231
7.8	Integration in verschiedenen Koordinatensystemen	233

7.8.1	Polarkoordinaten	233
7.8.2	Integration in Polarkoordinaten	234
7.9	*Integration nach Koordinatentransformationen	236
7.9.1	*Integration in Kugelkoordinaten	237
7.10	Kurzkurs Optimierung im $\mathbb{R}^n$	240
7.11	Vektoranalysis	241
7.11.1	Vektorielltes Kurvenintegral und Potential	243
7.11.2	Quellen und Senken von Vektorfeldern	247
7.11.3	Flächenelemente von Oberflächen	252
7.11.4	Der Gauß'sche Integralsatz	253
7.11.5	Satz von Stokes	254
8	Gewöhnliche Differentialgleichungen und Dynamische Systeme	255
8.1	Systeme mit einer Zustandsvariablen	264
8.2	Der harmonische Oszillator	266
8.2.1	Lösungsansatz im Reellen	267
8.2.2	Lösungsansatz im Komplexen	268
8.2.3	Der gedämpfte harmonische Oszillator	270
8.2.4	Lösungsansatz im Komplexen	271
8.3	Lineare dynamische Systeme	272
8.3.1	Stabilität und Eigenwerte	277
8.4	Nichtlineare autonome Systeme	279
8.4.1	Fixpunkte und Stabilität	280
8.4.2	Bifurkationen	286
8.4.3	Numerische Lösungsverfahren	291
8.5	Zeitdiskrete dynamische Systeme	293
8.5.1	Lineare Systeme	294
8.5.2	Nichtlineare Systeme	297
9	Wahrscheinlichkeitstheorie	299
9.1	Endliche Wahrscheinlichkeitsräume	299
9.1.1	Elementare Definitionen	300
9.1.2	Bedingte Wahrscheinlichkeit	303
9.1.3	Unabhängigkeit von Ereignissen	310
9.1.4	Produktexperimente	311
9.1.5	Zufallsvariablen	315
9.1.6	Erwartungswert, Varianz, Kovarianz	318
9.1.7	Das schwache Gesetz der großen Zahlen	328
9.2	Unendliche Wahrscheinlichkeitsräume	330
9.2.1	Diskrete Wahrscheinlichkeitsräume	330
9.2.2	Kontinuierliche Wahrscheinlichkeitsräume	332

10 Statistik	345
10.1 Parameterschätzung	345
10.1.1 Schätzprobleme und Schätzer	346
10.1.2 Eigenschaften von Schätzern	347
10.1.3 Konfidenzintervalle	352
10.1.4 Empirischer Median einer Stichprobe	355
10.2 Hypothesentest	356
10.2.1 Hilfsmittel	356
10.2.2 Ablehnungs- und Verträglichkeitsbereich	357
10.2.3 Der t -Test	359
10.2.4 Test auf Häufigkeiten	360
10.2.5 Test auf Einhaltung eines Grenzwerts	361

Einführung

Wozu brauchen Sie als angehende Naturwissenschaftlerin oder angehender Naturwissenschaftler Mathematik? Wir nennen Ihnen hier folgende Gründe:

- Die Mathematik stellt die **Sprache** für die Naturwissenschaften bereit und erlaubt somit, viele Sachverhalte überhaupt erst richtig zu formulieren. Sie ist also notwendige Basis zum Verständnis nicht nur von Physik und Chemie, sondern zunehmend auch der gesamten Biowissenschaften.
- In der Vorlesung werden einige **mathematische Verfahren**, z.B. statistische Tests, vorgestellt, die Sie für andere Vorlesungen, Praktika etc. benötigen. Dabei soll Ihr auch auf andere Methoden übertragbares Verständnis im Vordergrund stehen. Eine auch nur annähernd vollständige Behandlung aller Methoden oder intensives Einüben spezieller Berechnungen sind nämlich nicht möglich. Wichtig ist vielmehr, dass Sie sich bei Bedarf neue Verfahren auch selbständig aneignen können.
- Dennoch sollen Sie eine gewisse Fingerfertigkeit im Umgang mit den grundlegenden und wichtigsten **Rechenmethoden** erlangen, wie z.B. Differenzieren, Integrieren, Matrizenrechnung sowie Umformung von mathematischen Gleichungen, um einfache Anwendungen davon bei der Lektüre von Fachliteratur schnell nachvollziehen oder selber durchführen zu können.
- Darüber hinaus soll eine „**mathematische Denkweise**“ vermittelt oder vertieft werden, eine Fähigkeit zur Abstraktion, Analyse von Sachverhalten, Problemen etc., zur rationalen Argumentation und kritisch prüfender Sichtweise. Im Idealfall profitieren Sie davon auch in allen anderen wissenschaftlichen Disziplinen.
- Die Entwicklung der Computertechnik bietet großartige Möglichkeiten, mit Hilfe **mathematischer Modelle** nicht nur Vorhersagen zu treffen, sondern auch Parameter zu schätzen, Prozesse zu optimieren, Experimente besser zu planen etc. Ein moderner interdisziplinärer Wissenschaftszweig, das sog. **Wissenschaftliche Rechnen** beschäftigt sich mit dieser Thematik. Ein wichtiges Ziel unseres Mathematik-Kurses ist deshalb, Sie in die Lage zu versetzen, selbst mathematische Modelle zu verstehen, zu entwickeln und damit auf dem **Computer** zu arbeiten. Auch dafür ist es wichtig, die mathematischen Sprechweisen zu kennen, nicht zuletzt, um später auch mit Mathematikern oder mathematisch denkenden Naturwissenschaftlern effizient zusammenarbeiten zu können.

Zu Beginn des Kurses behandeln wir in etwa die gleichen Dinge, die auch in den Grundvorlesungen für Physiker oder Mathematiker behandelt werden – sie sind die Grundlage für fast alle Anwendungen der Mathematik. Allerdings werden wir wesentlich weniger Beweise durchführen, und mehr Wert auf praktische Rechenbeispiele legen. Ein Vorteil davon, sich an den mathematischen Grundvorlesungen zu orientieren, ist, dass Sie von Anfang an an die Denk- und Sprechweise der Mathematiker gewöhnt werden und viele der Begriffe lernen, die *jedem* mathematisch orientierten Wissenschaftler, also auch Physikern, Ingenieuren, Informatikern etc. geläufig sind. Dies wird Ihnen später die Kommunikation mit diesen Fachleuten erleichtern.

Der spätere Kursinhalt ist auf Ihren Bedarf als Naturwissenschaftler abgestimmt. Speziell werden wir uns auch mit mathematischer Modellierung und sogenannten *dynamischen Systemen* beschäftigen, um die Grundlage dafür zu schaffen, dass Sie später einmal eigenständig oder in interdisziplinärer Kooperation mathematische Modelle naturwissenschaftlicher Prozesse verstehen und entwickeln können.

Inhalt

1. Wir beginnen den Kurs mit einer **Einführung in die mathematische Logik**, und Sie erlernen gleich zu Beginn die Kurzsprache, in der vieles kürzer und genauer als mit Worten gesagt werden kann. Lassen Sie sich von den vielen neuen Symbolen nicht verwirren, Sie gewöhnen sich schnell daran.
2. Im Kapitel, **Lineare Algebra I**, befassen wir uns auf eine mathematische Weise mit dem Begriff des Raums und lernen wichtige Konzepte und Lösungsmethoden für sogenannte *lineare Gleichungssysteme* kennen, die häufig in mathematischen Anwendungen auftreten.
3. Im Kapitel **Lineare Algebra II** werden wir die Begriffe Determinante und Basistransformation behandeln, und sogenannte *Eigenwerte* von Matrizen kennenlernen, die für die Praxis so grundlegende Phänomene wie z.B. Resonanz oder Abklingverhalten beschreiben.
4. Das Kapitel **Analysis über $\mathbb{R}$** befasst sich mit Folgen und Reihen, der in der Praxis äußerst wichtigen Exponentialfunktion und der Logarithmus, sowie mit Ableitungen von Funktionen und Taylorentwicklung, Begriffe, denen man in der mathematischen Praxis überall begegnet.
Zum Abschluss des Kapitels Analysis werden wir die Integration als Umkehrung der Ableitung mathematisch exakt definieren und wichtige Rechentechniken behandeln. Diese werden Ihnen insbesondere in Physik und theoretischer Chemie von Nutzen sein.
5. Im Kapitel **Komplexe Zahlen** werden wir uns mit den komplexen Zahlen vertraut machen, die heutzutage zum unentbehrlichen Handwerkszeug vieler Praktiker gehören.
6. Das Kapitel **Skalarprodukte und Orthogonalität**, das zum Bereich der Linearen Algebra gehört, wirft ein neues Licht auf die Geometrie des Raumes. Eine sehr weitreichende Erkenntnis ist, dass auch Funktionen als Vektoren behandelt werden können, und man Begriffe

wie Abstand oder Winkel sinnvoll auf Funktionen verallgemeinern kann. Dies ermöglicht Techniken wie die Fourierzerlegung zu verstehen, die wie das menschliche Ohr aus einem Signal einzelne Frequenzen herausfiltert. Ausserdem führen wir den Begriff „selbstadjungierter Operator“ ein, der Ihnen in der theoretischen Chemie häufig begegnen wird.

7. In der **Analysis im $\mathbb{R}^n$** , die wieder ein mathematisches Grundlagenkapitel darstellt, werden wir Ableitungen und Integrale von Funktionen mehrerer Argumente behandeln.
8. Im letzten Kapitel über **Dynamische Systeme** behandeln wir die sogenannten „gewöhnlichen Differentialgleichungen“, mit deren Hilfe eine Vielzahl von Prozessen in der Biotechnologie modelliert werden kann. Mit Hilfe dieser Modelle kann man Vorhersagen treffen, Parameter schätzen, Hypothesen testen, oder sogar Prozesse mit Hilfe des Computers optimieren.
9. In der **Wahrscheinlichkeitstheorie** führen wir Begriffe wie Zufallsexperiment, bedingte Wahrscheinlichkeit (Formel von Bayes), Erwartungswert, Streuung und Korrelation ein, die insbesondere eine Grundlage für die Statistik bilden.
10. Die **Statistik** behandeln wir relativ ausführlich für einen mathematischen Grundkurs, denn Sie benötigen sie für die Planung, Auswertung und Interpretation fast aller Experimente und experimentellen Studien, die sie später durchführen werden.

Motivation, Ziele der Vorlesung und Bedeutung der Mathematik

- Die mathematische Denkweise und ein problemlösendes Denken zu schulen.
- Die exakte Formulierung von naturwissenschaftlichen Problemen soll erlernt werden: Mathematik ist die formale Sprache der Naturwissenschaften und das „Auge“ für komplexe Zusammenhänge.
- Es sollen technische Hilfsmittel für mathematische Berechnungen bereitgestellt werden.
- Eine Grundlage für den Einsatz von Computern und das Gebiet des Wissenschaftlichen Rechnens soll geschaffen werden.
- Ziel ist auch, die Fähigkeit zu Abstraktionen und das Erkennen von allgemeinen und tiefliegenden wissenschaftlichen Zusammenhängen zu erlernen.
- Speziell wird auch das mathematische Rüstzeug für das Verständnis von Physik, Physikalischer und Theoretischer Chemie vermittelt.
- Schließlich soll die Fähigkeit zur Kommunikation mit Mathematikern geschult werden (Interdisziplinarität). Diese Fähigkeit ist heutzutage im wissenschaftlichen und ausserwissenschaftlichen Berufsleben eine Schlüsselqualifikation.

Einige Tipps zum Lernen

- Das Verständnis von Konzepten, Ideen und Prinzipien ist viel wichtiger als stures Auswendiglernen von Details und technischen Arbeitsschritten !
- Mathematik lernt man nur, indem man sie betreibt und anwendet, zuhören und Stoff lernen reichen **nicht** (!)
- In der Vorlesung nicht unbedingt alles mitschreiben, der Stoff steht ja im Skript, sondern zuhören und verstehen !
- Mathematik kann durch ihre Strenge der Logik und tiefen Zusammenhänge faszinieren und inspirieren. Lassen Sie sich darauf ein und nicht abschrecken von anfänglich abstrakt und komplex anmutenden Konzepten.

Bedeutung der Mathematik in den modernen Biowissenschaften

- Biologie war in den letzten Jahrzehnten weitgehend eine reduktionistische Beschreibung einzelner Komponenten zellulärer Systeme auf molekularer Ebene
- in Zukunft liegt der Fokus auf:
 - quantitativen Studien
 - Systematischer Beschreibung von Zusammenhängen zwischen Komponenten, Dynamik und Funktionalität (Post-Genom-Ära).
→ zentrale Rolle der Mathematik

Tipps zum Lesen des Skriptes

Mathematisches Verständnis kommt eher in Form von plötzlichen Aha-Erlebnissen als durch stures Einpauken (abgesehen von einigen Rechentechniken, die einfach auch Training erfordern). Deshalb empfehlen wir Passagen, die für Sie schwer verständlich sind, zunächst einfach querzulesen und sich nicht gleich darin festzuhaken. Stattdessen kann man erst einmal versuchen, woanders Hilfe zu finden, z.B. im Gespräch mit Kommilitonen oder in anderen Büchern, und manchmal geht es dann ganz leicht; oder man liest einfach weiter und hofft, dass einem in einer späteren Textpartie doch noch alles klar wird. Danach kann und sollte man den schwierigen Textteil noch einmal lesen, oft geht es dann schon viel einfacher.

Universitätsbibliothek: Wir möchten Ihnen den Tipp geben, gleich zu Beginn der Vorlesung einmal in die Universitätsbibliothek zu gehen und auf jeden Fall neben dem Skript auch in andere Lehrbücher reinzuschauen, denn jeder hat andere Bedürfnisse und einen anderen Geschmack: oft versteht man mathematische Sachverhalte ganz augenblicklich, sobald man die für sich richtige Erklärung in irgendeinem Buch gefunden hat. Im nächsten Abschnitt geben wir einige Literaturempfehlungen.

Stift und Papier: Es empfiehlt sich außerdem beim Lesen mathematischer Texte, immer einen Stift und einen Haufen Papier zur Hand zu haben, auf dem man sich Dinge skizzieren oder Zwischenrechnungen durchführen kann. Dabei muss man überhaupt nicht den Anspruch haben, dass die beschriebenen Zettel am Ende schön aussehen; alles, was dem Verständnis dient, ist erlaubt! Sobald man etwas verstanden hat, kann man die meisten Zettel ja auch einfach wegwerfen, und nur die behalten, von denen man glaubt, dass Sie einem beim nächsten Lesen weiterhelfen. Es hilft dann, sich die entsprechenden Seitenzahlen im Skript auf die Zettel zu schreiben.

Sternchen: Da wir sehr viel Stoff in kurzer Zeit durchnehmen, können wir manche Gebiete nur sehr oberflächlich behandeln. Um Ihnen aber die Chance zu geben, einige für die Mathematik wichtige Begriffe kennenzulernen, die wir aber aus Zeitmangel hier nicht detailliert behandeln, haben wir viele Bemerkungen, Sätze, Abschnitte etc. hinzugefügt, die mit einem Sternchen (*) markiert sind, und die nicht unbedingt notwendig für das Verständnis des Kurses sind. Sie erlauben Ihnen, wenn Sie noch etwas weiter gehendes Interesse an einem Gebiet haben, noch etwas mehr dazuzulernen, das wir für interessant halten.

Index: Auf der Suche nach einem Stichwort kann man den Index verwenden. Die Wörter im Index haben oft mehrere Seitenangaben; zur Hervorhebung haben wir die wichtigste dieser Seitenzahlen jeweils fett gedruckt. Verweise auf Abbildungen sind kursiv.

Hyperlinks: Wir möchten Sie außerdem darauf hinweisen, dass dieses Skript in seiner elektronischen Version (als PDF mit Acrobat Reader geöffnet) ein Hypertext ist. Das heißt, dass Sie Querverweisen auf Definitionen, Formeln etc. durch einfachen Mausklick folgen können. Dies ist insbesondere beim Nachschlagen im Index oder im Inhaltsverzeichnis sehr praktisch.

Fehler und Feedback: Als letztes wollen wir noch den Hinweis geben, dass das vorliegende Skript trotz sorgfältigen Fehlerlesens sicher noch viele Fehler und Inkonsistenzen enthält. Auch wenn wir behaupten könnten, dies sei reine Absicht und diene nur dazu, Ihr kritisches Urteilsvermögen wachzuhalten, haben wir das Ziel, das Skript möglichst fehlerfrei werden zu lassen. Deshalb bitten wir Sie, wenn immer Sie beim Lesen des Skripts Fehler finden, oder wenn Sie sonstige Verbesserungsvorschläge haben, sich gleich beim Lesen die Seitenzahl und Ihren Änderungswunsch zu notieren. Und senden Sie uns bitte nach Sammeln Ihrer Korrekturen eine kleine Email an dirk.lebiedz@biologie.uni-freiburg.de mit dem Betreff „Korrektur zum Mathe Skript“.

Literaturempfehlungen

Zur Begleitung der Vorlesung, zum Vertiefen des Stoffes und zum Nacharbeiten, möchten wir Ihnen hier direkt einige Bücher empfehlen, die sie fast alle in der Uni-Bibliothek ausleihen können. Allgemeine Bücher, die das Thema Mathematik für Naturwissenschaftler behandeln, sind

- „Einführung in die Mathematik für Biologen“ von Eduard Batschelet [Bat80], das sehr viele schöne Beispiele enthält und auch die grundlegendsten Rechentechniken noch einmal behandelt, und
- „Grundkurs Mathematik für Biologen“ von Herbert Vogt [Vog94], das in kompakter Form die wichtigsten Konzepte behandelt und besonders die im zweiten Semester wichtige Statistik ausführlich behandelt.
- „Mathematik für Ingenieure und Naturwissenschaftler“ von Lothar Papula [Pap].
- Eher physikalisch interessierten Lesern gefällt vielleicht auch das Buch „Mathematik für Physiker“ von Fischer und Kaul [FK90].
- Ein kompaktes Nachschlagewerk und beliebtes Hilfsmittel für alle mathematisch arbeitenden Naturwissenschaftler ist das „Taschenbuch der Mathematik“ von Bronstein et al. [BSMM00].

Zur Nacharbeitung des Stoffes in Analysis empfehlen wir Ihnen eines oder mehrere der folgenden Lehrbücher:

- „Analysis I“ von Forster [Fora], das schön kompakt, aber auch sehr abstrakt ist und sich an Mathematikstudenten wendet.
- „Folgen und Funktionen: Einführung in die Analysis“ von Harald Scheid [Sch], das viele Beispiele enthält und ursprünglich für Lehramtsstudenten gedacht war.
- „Analysis I“ von Martin Barner und Friedrich Flohr [BF].
- „Calculus“ von S. L. Salas und Einar Hille [SH], das viele Erläuterungen und sehr ausführliche Beispiele enthält.
- „Analysis I“ von H. Amann und J. Escher [AE99].

Zum Themengebiet der Linearen Algebra empfehlen wir Ihnen die folgenden Lehrbücher:

- „Lineare Algebra“ von Klaus Jähnich [Jäh98], ein Buch mit vielen graphischen Veranschaulichungen, das wir zur Vertiefung und Nacharbeitung des Stoffes in Linearer Algebra empfehlen.
- „Lineare Algebra. Schaum’s Überblicke und Aufgaben“ von Seymour Lipschutz [Lip99], das auch gut zur Nacharbeitung des Stoffes in Linearer Algebra geeignet ist und viele schöne Beispiele enthält und alles schön ausführlich erklärt.
- „Lineare Algebra“ von Gerd Fischer [Fis00], das wie „Analysis I“ von Forster schön kompakt ist, aber sich primär an Mathematikstudenten wendet.
- „Übungsbuch zur Linearen Algebra“ von H. Stoppel and B. Griesse [SG], wenn man zum besseren Verständnis noch extra Übungsaufgaben sucht.

Zur Wahrscheinlichkeitstheorie und Statistik können wir die folgenden Lehrbücher empfehlen:

- Ulrich Krengel, „Einführung in die Wahrscheinlichkeitstheorie und Statistik“ [Kre02], das wir als Vorlage zur Konzipierung dieser Vorlesung benutzt haben.
- Karl Bosch: „Elementare Einführung in die Wahrscheinlichkeitsrechnung“ [Bos99] und „Elementare Einführung in die angewandte Statistik“ [Bos00].
- „Angewandte Statistik“ von Lothar Sachs [Sac02], mit vielen ausführlich dargestellten Beispielen.
- „Statistische Datenanalyse“ von Werner A. Stahel [Sta02].
- sowie das auf biologische Anwendungen ausgerichtete Standardwerk „Biometry“ von Sokal und Rohlf [SR94].
- und außerdem die unterhaltsamen wie informativen populärwissenschaftlichen Bücher [BBDH01] und [Krä00], die viele grundlegende Ideen der Wahrscheinlichkeitstheorie und Statistik illustrieren und insbesondere vor dem falschen Gebrauch von Statistik warnen.

Für den Bereich Dynamische Systeme und Modellierung in der Biologie empfehlen wir:

- Strogatz: „Nonlinear Dynamics and Chaos“ [Str00].
- Walter: „Gewöhnliche Differentialgleichungen“ [Wal93]. Ein an Mathematiker gerichtetes Einführungswerk.
- Ebenso an Mathematiker wendet sich das Buch „Gewöhnliche Differentialgleichungen“ von Amann [Ama83].
- „Analysis II“ von Forster [Forb], das schön kompakt, aber auch sehr abstrakt ist und sich an Mathematikstudenten wendet.
- Das Buch „Modeling Dynamic Phenomena in Molecular and Cellular Biology“ von L. Segel [Seg84] enthält und diskutiert viele interessante Modelle aus der molekularen Biologie.
- Yeagers et al.: „An Introduction to the Mathematics of Biology, With Computer Algebra Models.“ [YHYS96]. Dieses Buch enthält viele Computermodelle in der auch bei Biologen populären Software MATHEMATICA.
- „Mathematical Biology“ von J.D. Murray [Mur02] ist eine wunderbare und ausführliche Sammlung von mathematischen Modellen in der Biologie.

Anmerkungen zur Entstehung dieses Skriptes

Dieses Skript entstand aus einer zweisemestrigen Vorlesung *Mathematik* mit insgesamt 8 SWS Vorlesungen und 4 SWS Übungen, die die Autoren in den akademischen Jahren 2002/2003 (M. Diehl und T. Fischer) bzw. 2006/2007 (D. Lebiedz) an der Universität Heidelberg im Rahmen des Bachelor of Science Studiengangs *Molekulare Biotechnologie* gehalten haben. Die Niederschrift des vorliegenden Skriptes wurde von vier als wissenschaftliche Hilfskräfte/Mitarbeiter beschäftigten Studenten Lorenz Steinbock, Kristian Wadel, Melanie Franzem und Kai Antweiler tatkräftig unterstützt.

Kapitel 1

Einführung in die mathematische Logik

Die gewöhnliche Alltagssprache kann formalisiert werden. Dies erlaubt, mit klar definierten Symbolen auch komplexe Sachverhalte so auszudrücken, dass sie jeder Mensch, der die mathematische Symbolsprache kennt, auf genau die gleiche Weise versteht. Ein glücklicher Umstand ist die Tatsache, dass die mathematische Symbolsprache international verstanden wird: man kann die gleichen Symbole in Indien ebenso wie in Algerien, in Japan ebenso wie in Argentinien verwenden.

1.1 Aussagen und logische Verknüpfungen

Im Zentrum der mathematischen Logik stehen **Aussagen**, wie z.B. „Es ist kalt“ oder „ $2+2=5$ “. Mit dem Symbol $:\Leftrightarrow$ kann man einer Aussagenvariable A einen Aussagen-Wert wie z.B. „Es ist kalt“ zuweisen:

$$A :\Leftrightarrow \text{„Es ist kalt“}, \quad \text{oder} \quad B :\Leftrightarrow \text{„Ich friere“},$$

ganz analog wie man z.B. einer Zahl-Variable a den Wert $a := 3$ zuweisen kann. Man kann das Symbol $:\Leftrightarrow$ als „wird definiert als“ oder „ist per Definition äquivalent zu“ lesen. Wir sammeln nun einige wichtige Tatsachen über Aussagen.

- Aussagen in der Mathematik sind entweder **wahr** oder **falsch**; man sagt, sie haben den **Wahrheitswert** w oder f (Engl.: true/false). Erstaunlicherweise sind sich Mathematiker nahezu immer einig, ob eine Aussage wahr oder falsch ist, z.B. ist „ $2+2=5$ “ falsch, aber „ $2+2=4$ “ wahr.
- Aussagen, die den gleichen Wahrheitswert haben, heissen **äquivalent**. Sind zwei Aussagen A und B äquivalent, schreibt man $A \Leftrightarrow B$. Man spricht dies auch als „A genau dann, wenn B“ oder sogar „A dann und nur dann, wenn B“ (Engl.: „if and only if“, kurz auch manchmal geschrieben als „iff“). Die Äquivalenz ist sozusagen die Gleichheit von Aussagen. Ein Beispiel dafür hatten wir ja schon in dem Symbol $:\Leftrightarrow$ kennengelernt, das einfach *definiert*, dass zwei Aussagen äquivalent (gleich) sein sollen. Ein weiteres Beispiel ist die folgende

Äquivalenz¹:

$$(a = 5) \Leftrightarrow (2a = 10),$$

denn ganz egal welchen Wert die Zahlvariable a hat, ist jede der beiden Aussagen genau dann wahr, wenn die andere wahr ist.

- Aussagen A können **verneint** werden, und werden dadurch zu einer neuen Aussage, der **Negation** von A , dargestellt durch das Symbol $\neg A$. Man liest dies auch als „Aussage A ist falsch.“ Z.B. gilt

$$\neg(\text{„Mir ist kalt.“}) \Leftrightarrow \text{„Mir ist nicht kalt.“}$$

oder auch

$$\neg(2 + 2 = 5) \Leftrightarrow (2 + 2 \neq 5)$$

- Die doppelte Verneinung neutralisiert die einfache Verneinung, genau wie in der gesprochenen Sprache:

$$\neg(\neg A) \Leftrightarrow A \quad (\text{„Es ist falsch, dass } A \text{ falsch ist.“})$$

- Zwei Aussagen A und B können durch die UND-Verknüpfung (Konjunktion) zu einer neuen Aussage verknüpft werden :

$$A \wedge B :\Leftrightarrow \text{„}A \text{ und } B\text{“},$$

z.B. $A \wedge B \Leftrightarrow$ „Es ist kalt **und** ich friere“

Diese Aussage ist nur dann wahr, wenn A und B beide wahr sind.

- Eine andere Verknüpfung ist die ODER-Verknüpfung (Disjunktion):

$$A \vee B :\Leftrightarrow \text{„}A \text{ oder } B\text{“}.$$

Die Aussage $A \vee B$ ist wahr, wenn A oder B wahr sind, oder wenn beide zugleich wahr sind.

Achtung: Das mathematische „oder“ ist ein **einschliessendes oder**, kein „entweder-oder“. Beispiel: $A \vee B \Leftrightarrow$ „Es ist kalt **und/oder** ich friere.“

- Man kann logische Verknüpfungen wie z.B. die UND- oder die ODER- Verknüpfung auch über eine sogenannte **Wahrheitstafel** repräsentieren, in die man alle möglichen Kombinationen von Wahrheitswerten, die A und B annehmen können, in die ersten beiden Spalten schreibt, und dann die Ergebnis-Werte, die die Verknüpfungen haben, in die folgenden Spalten:

A	B	$A \wedge B$	$A \vee B$
w	w	w	w
w	f	f	w
f	w	f	w
f	f	f	f

¹Strenggenommen ist $(a = 5)$ nur dann eine Aussage, wenn a einen festen Wert hat. Sonst ist es eine sogenannte Aussageform, die wir aber erst in Abschnitt 1.2 einführen werden.

Man kann auch Wahrheitstafeln für Negation und Äquivalenz aufstellen:

A	$\neg A$	und	A	B	$A \Leftrightarrow B$
w	f		w	w	w
f	w		w	f	f
			f	w	f
			f	f	w

- Mit Hilfe von „ $\neg$ “, „ $\wedge$ “, „ $\vee$ “ kann *jede* mögliche Verknüpfung hergestellt werden. Als ein Beispiel betrachten wir z.B. die „entweder-oder“ Verknüpfung. Man kann „Entweder A oder B “ tatsächlich darstellen als

$$(A \wedge (\neg B)) \vee ((\neg A) \wedge B),$$

wie wir anhand der Wahrheitstafeln überprüfen können:

A	B	$\neg A$	$\neg B$	$A \wedge (\neg B)$	$(\neg A) \wedge B$	$(A \wedge (\neg B)) \vee ((\neg A) \wedge B)$
w	w	f	f	f	f	f
w	f	f	w	w	f	w
f	w	w	f	f	w	w
f	f	w	w	f	f	f

Die letzte Spalte entspricht tatsächlich der gewünschten Wahrheitstafel von „Entweder A oder B “.

Für Interessierte: Man kann nur aus „ $\neg$ “, „ $\vee$ “ allein alle anderen Verknüpfungen aufbauen. Wie erzeugt man aus diesen beiden z.B. „ $\wedge$ “? Es geht sogar noch kompakter, und im Prinzip reicht sogar nur eine einzige Verknüpfung, nämlich „Weder-A-noch-B“, um alle anderen daraus aufzubauen. Wie macht man daraus „ $\neg$ “ und „ $\vee$ “?

- Man kann leicht mit der Wahrheitstafel den **Satz von De Morgan** zeigen:

$$\neg(A \wedge B) \Leftrightarrow (\neg A) \vee (\neg B)$$

und

$$\neg(A \vee B) \Leftrightarrow (\neg A) \wedge (\neg B).$$

Illustration: „Es ist falsch, dass es kalt ist und ich friere“ ist das gleiche wie „Es ist nicht kalt und/oder ich friere nicht“

- Interessant ist die Definition der sogenannten **Implikation**

$$A \Rightarrow B :\Leftrightarrow \text{„Aus } A \text{ folgt } B\text{“}$$

Die Aussage $A \Rightarrow B$ ist sicher falsch, wenn A richtig und B falsch ist. Man definiert nun einfach, dass sie sonst immer wahr ist. Diese Definition macht Sinn, wie wir bald sehen werden. Die Wahrheitstafel hat also die Form:

A	B	$A \Rightarrow B$
w	w	w
w	f	f
f	w	w
f	f	w

$A \Rightarrow B$ ist übrigens äquivalent zur Aussage $(\neg A) \vee B$, wie man anhand der Wahrheitstafel nachprüfen kann. Interessant ist auch, dass die Äquivalenz $A \Leftrightarrow B$ selbst äquivalent zur Aussage $(A \Rightarrow B) \wedge (B \Rightarrow A)$ ist.

- Falls eine Aussage der Form $(A \Rightarrow B) \wedge (B \Rightarrow C)$ (kurz: $A \Rightarrow B \Rightarrow C$) gilt, so ist A eine **hinreichende** Bedingung für B , denn sie reicht aus, um die Wahrheit von B zu folgern. Andererseits ist C eine **notwendige** Bedingung für B , denn wenn B wahr sein soll, so ist C notwendig auch wahr. Man kann sich dies gut anhand der hinreichenden und notwendigen Bedingungen, wann ein Punkt x ein Minimum einer Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ ist, merken, die vielen aus der Schule bekannt sind: Es gilt nämlich für alle $x \in \mathbb{R}$

$$\left(f'(x) = 0 \wedge f''(x) > 0 \right) \Rightarrow x \text{ ist Minimum von } f \Rightarrow f'(x) = 0.$$

1.2 Aussageformen und Quantoren

Aussagen können auch von Variablen abhängen. Man spricht dann von einer **Aussageform**. Beispiele:

$$\begin{aligned} A(x) &: \Leftrightarrow \text{„Person } x \text{ hat ein Gehirn“} \\ B(x, y) &: \Leftrightarrow \text{„Person } x \text{ ist mit Person } y \text{ verheiratet“} \\ C(n) &: \Leftrightarrow \text{„Die Zahl } n \text{ ist durch 2 teilbar“} \\ D(a) &: \Leftrightarrow (a = 5) \end{aligned}$$

(wobei wir die letzte Aussageform schon früher verwendet haben). Eine Aussageform $A(\cdot)$ ist im strengen Sinne keine Aussage, denn erst wenn man einen bestimmten Wert in die Variable x einsetzt, hat sie einen bestimmten Wahrheitswert und wird zu einer bestimmten Aussage, nämlich zu $A(x)$.

- Die Variablen können nur Werte aus bestimmten **Mengen** annehmen, z.B.

$$\begin{aligned} \mathbb{X} &:= \text{„Menge aller Personen im Hörsaal“} = \{\text{Michael, Severine, } \dots\}, \\ \mathbb{N} &:= \text{„Menge der natürlichen Zahlen“} = \{0, 1, 2, 3, \dots\}, \\ \mathbb{R} &:= \text{„Menge der reellen Zahlen“}. \end{aligned}$$

Die Aussageform $C(n)$ = „Die Zahl n ist durch 2 teilbar“ nimmt z.B. für jeden Wert $n \in \mathbb{N}$ einen Wahrheitswert an, und wird damit zu einer Aussage (z.B. ist $C(4)$ wahr und $C(5)$ falsch).

- Aussageformen können verwendet werden, um neue Mengen zu definieren. Die Menge aller Elemente x aus $\mathbb{X}$, für die die Aussage $A(x)$ wahr ist, bezeichnet man mit

$$\{x \in \mathbb{X} | A(x)\}.$$

In unserem Beispiel wäre dies also die Menge aller Personen im Hörsaal, die ein Gehirn haben. Ein anderes Beispiel wäre die Menge aller positiven reellen Zahlen:

$$\mathbb{R}_+ := \{x \in \mathbb{R} | x > 0\}.$$

Eine wichtige Möglichkeit, aus Aussageformen Aussagen zu machen, sind Aussagen der Art: „Alle Personen im Hörsaal haben ein Gehirn“ oder „Mindestens eine Person im Hörsaal hat ein Gehirn“. In der mathematischen Symbolsprache erfolgt dies mit Hilfe von sogenannten **Quantoren**:

- Man benutzt den Allquantor „ $\forall$ “ um zu sagen „für alle“, also z.B.

$$\forall x \in \mathbb{X} : A(x) :\Leftrightarrow \text{„Für alle } x \text{ aus } \mathbb{X} \text{ gilt: } A(x)\text{“}$$

Mit den oben stehenden Definitionen von $\mathbb{X}$ und $A(x)$ hieße dies also: „Für jede Person x im Hörsaal gilt, dass x ein Gehirn hat.“

- und den Existenzquantor „ $\exists$ “ um zu sagen „es existiert mindestens ein“, also z.B.

$$\exists x \in \mathbb{X} : A(x) :\Leftrightarrow \text{„Es existiert mindestens ein } x \text{ aus } \mathbb{X} \text{ für das gilt: } A(x)\text{“}$$

Dies hieße also „Es gibt mindestens eine Person x im Hörsaal, so dass x ein Gehirn hat.“

- Sind nicht alle Variablen einer Aussageform durch Quantoren quantifiziert, bleibt eine neue Aussageform übrig. Mit obenstehender Definition von $B(x, y)$ und der Menge $\mathbb{Y}$ aller Menschen können wir z.B. eine Aussageform $E(x)$ definieren:

$$E(x) :\Leftrightarrow (\exists y \in \mathbb{Y} : B(x, y)),$$

also „Es gibt mindestens einen Menschen y , so dass Person x mit y verheiratet ist“ oder kurz „Person x ist verheiratet“.

- Man kann natürlich auch geschachtelte Aussagen durch doppelte Anwendung von Quantoren erzeugen, z.B.

$$\forall x \in \mathbb{X} : (\exists y \in \mathbb{Y} : B(x, y))$$

was man meist ohne Klammern als

$$\forall x \in \mathbb{X} \quad \exists y \in \mathbb{Y} : B(x, y)$$

schreibt, und was man liest als: „Für jedes x aus $\mathbb{X}$ gibt es ein y aus $\mathbb{Y}$ so dass $B(x, y)$ gilt.“ Im Beispiel wäre dies die Aussage „Für jede Person im Hörsaal gibt es (mindestens) einen Menschen, mit dem sie verheiratet ist.“ oder kurz „Alle Personen im Hörsaal sind verheiratet.“

- Die Verneinung von Aussagen oder Aussageformen, die Quantoren enthalten, folgt der Logik unserer Sprache: „Es ist falsch, dass für alle x die Aussage $A(x)$ gilt“ ist äquivalent zu „Es gibt mindestens ein x , so dass $A(x)$ nicht gilt“. Umgekehrt ist „Es ist falsch, dass es ein x mit $A(x)$ gibt“ äquivalent zu „Für kein x gilt $A(x)$ “. In Symbolschreibweise setzt man also:

$$\begin{aligned}\neg(\forall x \in \mathbb{X} : A(x)) &: \Leftrightarrow (\exists x \in \mathbb{X} : \neg A(x)) \quad \text{und} \\ \neg(\exists x \in \mathbb{X} : A(x)) &: \Leftrightarrow (\forall x \in \mathbb{X} : \neg A(x)).\end{aligned}$$

Mit dieser Definition kann man durch doppelte Anwendung auch geschachtelte Aussagen verneinen:

$$\begin{aligned}\neg\left(\forall x \in \mathbb{X} \quad \exists y \in \mathbb{Y} : B(x, y)\right) &\Leftrightarrow \left(\exists x \in \mathbb{X} \quad \forall y \in \mathbb{Y} : \neg B(x, y)\right) \\ \neg\left(\exists x \in \mathbb{X} \quad \forall y \in \mathbb{Y} : B(x, y)\right) &\Leftrightarrow \left(\forall x \in \mathbb{X} \quad \exists y \in \mathbb{Y} : \neg B(x, y)\right)\end{aligned}$$

Merkregel: „Beim Durchziehen der Verneinung von links nach rechts drehen sich alle Quantoren um.“

- Aussageformen können auch verknüpft werden. Die Aussageform „Wenn n durch 4 teilbar ist, dann ist n durch 2 teilbar“ kann z.B. aus den zwei Aussageformen $B(n) : \Leftrightarrow „n \text{ ist durch 4 teilbar}“$ und $C(n) : \Leftrightarrow „n \text{ ist durch 2 teilbar}“$ durch

$$A(n) : \Leftrightarrow (B(n) \Rightarrow C(n))$$

erhalten werden.

1.3 Wahre Aussagen in der Mathematik

Man könnte etwas überspitzt formulieren, dass das Ziel der Mathematik einfach nur ist, eine Menge von interessanten oder nützlichen Aussagen mit dem Wahrheitswert „wahr“ zu produzieren. Aber wie entscheidet man in der Mathematik, ob eine Aussage wahr ist? Ist z.B. die Aussage „Jede durch 4 teilbare Zahl ist auch durch 2 teilbar“ wahr oder falsch? Wir können diese Aussage in Symbolsprache ausdrücken, indem wir mit $B(n) := „n \text{ ist durch 4 teilbar}“$ und $C(n) := „n \text{ ist durch 2 teilbar}“$ schreiben:

$$A : \Leftrightarrow \left(\forall n \in \mathbb{N} : B(n) \Rightarrow C(n)\right).$$

Durch Einsetzen aller Werte n aus $\mathbb{N}$ und unter Verwendung der Wahrheitstafel der Implikation (die mit diesem Beispiel nachträglich gerechtfertigt wird), könnte man nun die komplette Wahrheitstafel erstellen, und erhielte:

n	$B(n)$	$C(n)$	$B(n) \Rightarrow C(n)$
0	w	w	w
1	f	f	w
2	f	w	w
3	f	f	w
4	w	w	w
5	f	f	w
$\vdots$	$\vdots$	$\vdots$	$\vdots$

Daraus könnte man vermuten, dass die Aussage wahr ist. Ein wirklicher Beweis mit dieser Methode würde allerdings unendlich lange dauern. Die Mathematiker haben sich deshalb für einen anderen Weg entschieden: sie beweisen die Gültigkeit einer Aussage, indem sie sich andere Aussagen zu Hilfe nehmen, deren Gültigkeit bereits anerkannt ist, und daraus die Wahrheit der betreffenden Aussage folgern.

- Die Mathematik startet mit **Definitionen**, die uns ja inzwischen wohlbekannt sind, und mit sogenannten **Axiomen**, das sind Aussagen, die per Definition als wahr gesetzt werden. Z.B. setzt man sich das Axiom: „Jede natürliche Zahl hat einen Nachfolger.“, mit dessen Hilfe man nun vieles andere beweisen kann.
- Eine Aussage, deren Wahrheit bewiesen wurde, heißt **Satz** oder **Theorem**. Sätze heißen manchmal auch **Lemma**, wenn sie als nicht so wichtig angesehen werden, oder auch **Korollar**, wenn sie aus einem anderen Satz sehr leicht gefolgert werden können.
- Eine Aussage, von der man ernsthaft glaubt, dass sie wahr ist, die aber noch nicht bewiesen ist, nennt man eine **Vermutung**. Z.B. wurde vom französischen Mathematiker Pierre de Fermat 1637 die sogenannte „Fermatsche Vermutung“ aufgestellt, die er als Randnotiz in seiner Ausgabe des antiken Buches „Arithmetica“ von Diophant schrieb:

$$\forall n, x, y, z \in \mathbb{N}, n \geq 3, x, y, z \geq 1 : x^n + y^n \neq z^n.$$

Fermat selbst behauptete zwar, er habe „hierfür einen wahrhaft wunderbaren Beweis, doch ist dieser Rand hier zu schmal, um ihn zu fassen“, aber das allein reichte natürlich nicht aus, um seiner Aussage den Status eines Satzes zu verleihen. Generationen von Mathematikern haben versucht, den Beweis „wiederzufinden“ (viele haben aber auch versucht, die Vermutung durch ein Gegenbeispiel zu widerlegen). Erst vor wenigen Jahren wurde sie von Andrew Wiles auf über 100 Seiten bewiesen (Annals of Mathematics, Mai 1995) und der Beweis wurde strengstens von anderen Mathematikern überprüft. Seitdem nennt man die obenstehende Aussage auch „Fermats letzten Satz“.

- Eine Aussage, von der man einfach einmal annimmt, dass sie wahr sei (ohne das ganz ernsthaft zu glauben), nennt man **Hypothese** oder auch **Annahme**. Dies hilft oft bei Beweisen, z.B. bei Fallunterscheidungen oder bei sog. Widerspruchsbeweisen.

- **Direkte Beweise** leiten einen Satz direkt aus anderen wahren Aussagen ab. Oft funktionieren Sie nach dem Muster: wenn $A \Rightarrow B$ und $B \Rightarrow C$ gilt, dann auch $A \Rightarrow C$, d.h. man geht Schritt für Schritt in Richtung der zu beweisenden Aussage.
- **Indirekte Beweise** oder **Widerspruchsbeweise** (auch **reductio ad absurdum**) nehmen zum Beweis einer Aussage A als zu widerlegende Hypothese einfach zunächst an, dass $\neg A$ wahr sei. Aus $\neg A$ leitet man dann auf direktem Wege eine eindeutig falsche Aussage her, und folgert daraus, dass $\neg A$ falsch, also A wahr ist.

1.4 Vollständige Induktion

1.4.1 Induktion und Deduktion

Im Duden Fremdwörterbuch wird **Induktion** als wissenschaftliche Methode beschrieben, bei der vom besonderen Einzelfall auf das Allgemeine, Gesetzmäßige geschlossen wird. Dies ist ein übliches Vorgehen in den Naturwissenschaften. Die Induktion hilft uns, Ideen für Gesetzmäßigkeiten zu bekommen. Ein großes Problem für die wahrheitsliebenden Mathematiker ist jedoch, dass die Gesetzmäßigkeit durch Induktion nur erraten wird, aber nicht bewiesen! Die Induktion steht damit im Gegensatz zur **Deduktion**, bei der eine Gesetzmäßigkeit aus bereits Bekanntem abgeleitet wird, und die eine völlig legitime Beweistechnik ist.

Zum Glück gibt es eine mathematisch korrekte Möglichkeit, vom Einzelfall auf das Allgemeine zu schließen, und diese Beweistechnik nennt sich **vollständige Induktion**. Es ist eine Technik, um Aussagen der Form

$$\forall n \in \mathbb{N} : A(n)$$

zu beweisen. Das Vorgehen illustrieren wir an einem Beispiel.

Beispiel 1.4.1. Wir betrachten die Zahlenfolge

$$1 + 3 + 5 + \cdots + (2n + 1) =: s_n. \quad (1.1)$$

Diese lässt sich auch durch folgende **Rekursionsformel** definieren.

$$s_0 = 1, \quad (1.2)$$

$$s_n = s_{n-1} + (2n + 1) \quad \text{für } n > 0. \quad (1.3)$$

Wir möchten eine explizite Formel für s_n finden, mit der wir s_n direkt berechnen können, ohne vorher $s_1, \dots, s_{n-1}$ ausrechnen oder, was auf das gleiche hinausläufe, $(n + 1)$ Zahlen summieren zu müssen.

Um eine solche Formel erraten zu können, berechnen wir s_n für die ersten paar n :

$$\begin{aligned} s_0 &= 1, \\ s_1 &= 1 + 3 = 4, \\ s_2 &= 4 + 5 = 9. \end{aligned}$$

Unsere naheliegende Vermutung ist, dass $(s_n)_{n \in \mathbb{N}}$ die Folge der Quadratzahlen ist. Diese Vermutung haben wir also mit Hilfe der normalen Induktion erhalten. Sie ist damit allerdings noch nicht bewiesen. Wir werden Sie sogleich mit Hilfe der **vollständigen** Induktion beweisen, und nennen Sie der Einfachheit halber jetzt bereits „Satz“.

Satz 1.4.2. Sei s_n durch (1.1) definiert. Dann gilt für alle $n \in \mathbb{N}$ die Aussage

$$A(n) : \Leftrightarrow (s_n = (n + 1)^2). \quad (1.4)$$

1.4.2 Technik der vollständigen Induktion

Die vollständigen Induktion geht zum Beweis der Aussage

$$\forall n \in \mathbb{N} : A(n)$$

folgendermaßen vor:

- 1) Wir zeigen zunächst, dass die Aussage $A(0)$ wahr ist. Dies nennt sich **Induktionsanfang**.
- 2) Dann zeigen wir im sogenannten **Induktionsschritt**, dass für jedes beliebige $n \in \mathbb{N}$ die Aussage $A(n + 1)$ wahr ist, wenn wir nur voraussetzen, dass $A(0), A(1), \dots, A(n)$ bereits wahr sind. Die für den Beweis benötigten Annahmen bezeichnet man als **Induktionsvoraussetzung**, die zu beweisende Aussage $A(n + 1)$ als **Induktionsbehauptung**. Man beweist also

$$\forall n \in \mathbb{N} : (A(0) \wedge A(1) \wedge \dots \wedge A(n)) \Rightarrow A(n + 1)$$

Wenn man sowohl Induktionsanfang als auch Induktionsschritt gemacht hat, kann man daraus sofort folgern, dass $A(n)$ für alle $n \in \mathbb{N}$ wahr ist.

Illustration am Beispiel 1.4.1

- 1) **Induktionsanfang:** Behauptung (1.4) ist für $n = 0$ wahr, denn

$$s_0 = 1 = (0 + 1)^2.$$

Damit ist $A(0)$ bereits bewiesen.

- 2) **Induktionsschritt:** Wir leiten aus der *Induktionsvoraussetzung* die *Induktionsbehauptung* her. In diesem Beispiel benötigen wir statt aller bereits bewiesenen Aussagen $A(0), A(1), \dots, A(n)$ nur die letzte, nämlich $A(n)$, als Voraussetzung.
Induktionsvoraussetzung: Sei Behauptung (1.4) für n wahr, also $s_n = (n + 1)^2$
Induktionsbehauptung: Behauptung (1.4) ist auch für $(n + 1)$ richtig.
Beweis der Induktionsbehauptung: Unter Verwendung der Rekursionsformel (1.3) und der Induktionsvoraussetzung erhalten wir

$$\begin{aligned}
s_{n+1} &= s_n + (2n + 3) \quad (\text{nach Rekursionsformel (1.3)}) \\
&= (n + 1)^2 + 2n + 3 \quad (\text{nach Induktionsvoraussetzung}) \\
&= (n + 1)^2 + 2(n + 1) + 1 \\
&= ((n + 1) + 1)^2 \\
&= (n + 2)^2.
\end{aligned}$$

Die Behauptung (1.4) ist also sowohl für $n = 0$ richtig und der Induktionsschritt ist bewiesen, somit gilt (1.4) nach dem Prinzip der vollständigen Induktion für alle $n \in \mathbb{N}$. $\square$

Bemerkung 1.4.3. Das Symbol $\square$ wird verwendet, um zu sagen, dass ein Beweis beendet ist. Wir bemerken noch, dass wir nicht zu allen im Skript angegebenen Sätzen einen Beweis liefern. Oft lassen wir einen solchen der Kürze halber weg. Bei einigen wichtigen Sätzen ist ein Beweis zu lang oder auch zu kompliziert und geht weit über das Niveau dieser Vorlesung hinaus.

Beispiel 1.4.4. Ein weiteres Beispiel für eine durch vollständige Induktion beweisbare Aussage ist die Bernoulli-Ungleichung.

Satz 1.4.5 (Bernoulli Ungleichung).

Sei $-1 \leq a \in \mathbb{R}$. Für alle $n \in \mathbb{N}$ mit $n \geq 1$ gilt

$$(1 + a)^n \geq 1 + na, \quad (1.5)$$

und die Gleichheit gilt nur für $n = 1$ oder $a = 0$.

Beweis: Da hier eine Behauptung für $\forall n \geq 1$ bewiesen werden soll, startet man hier nicht mit $n = 0$, sondern mit $n = 1$.

1) **Induktionsanfang:** Für $n = 1$ gilt

$$(1 + a)^1 = 1 + a = 1 + 1a.$$

2) **Induktionsschritt:** Seien die Behauptungen für n richtig. Dann gilt

$$\begin{aligned}
(1 + a)^{n+1} &= (1 + a)^n (1 + a) \\
&\geq (1 + na)(1 + a) \quad (\text{nach Induktionsvoraussetzung}) \\
&= 1 + (n + 1)a + na^2. \\
&\geq 1 + (n + 1)a \quad (\text{wegen } na^2 \geq 0).
\end{aligned} \quad (1.6)$$

Also gilt insgesamt $(1 + a)^{n+1} \geq 1 + (n + 1)a$. In (1.6) gilt in der zweiten Zeile (erste Ungleichung) Gleichheit genau dann, wenn $(1 + a)^n = 1 + na$, d.h., nach Induktionsvoraussetzung dann und nur dann, wenn $n = 1$ oder $a = 0$. In der vierten Zeile (zweite Ungleichung) gilt Gleichheit genau dann, wenn $a = 0$. Insgesamt gilt für $n \geq 2$ die Gleichheit also nur für $a = 0$. Damit sind alle Aussagen für den Induktionsschritt bewiesen. $\square$

1.5 Binomialkoeffizient und binomischer Lehrsatz

Am Ende dieses Kapitels über mathematische Logik möchten wir die gerade erlernte Methode der vollständigen Induktion gleich einmal anwenden, um den sogenannten Binomialkoeffizienten kennenzulernen, der insbesondere in der *Kombinatorik* eine große Rolle spielt, einem Teilgebiet der Mathematik, dass sich mit der „Zahl möglicher Anordnungen“ beschäftigt.

Zur Motivation des Binomialkoeffizienten entwickeln wir die Polynome $(x + y)^n$ für die ersten fünf natürlichen Exponenten n :

$$\begin{aligned}(x + y)^0 &= 1, \\(x + y)^1 &= x + y, \\(x + y)^2 &= x^2 + 2xy + y^2, \\(x + y)^3 &= x^3 + 3x^2y + 3xy^2 + y^3, \\(x + y)^4 &= x^4 + 4x^3y + 6x^2y^2 + 4xy^3 + y^4.\end{aligned}$$

Allgemein gilt:

Satz 1.5.1 (Binomischer Lehrsatz).

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

Für den Beweis durch vollständige Induktion verweisen wir auf die Lehrbücher, z.B. auf [Fora]. Dabei haben wir folgende Notation verwendet:

$$\binom{n}{k} := \begin{cases} \frac{n!}{(n-k)!k!} & \text{für } 0 \leq k \leq n \in \mathbb{N}, \\ 0 & \text{sonst,} \end{cases} \quad (1.7)$$

$$n! := \begin{cases} 1 & \text{für } n = 0, \\ \prod_{k=1}^n k & \text{für } 1 \leq n \in \mathbb{N}. \end{cases} \quad (1.8)$$

Den Ausdruck $n!$ lesen wir als „*n Fakultät*“ und den *Binomialkoeffizienten* $\binom{n}{k}$ als „*n über k*“. Die Binomialkoeffizienten ungleich Null, also mit $0 \leq k \leq n$, lassen sich im **Pascalschen Dreieck** anordnen (s. Abbildung 1.5.) In diesem erkennen wir das Muster der Koeffizienten in (1.7) wieder. Der Binomialkoeffizient $\binom{n}{k}$ steht im Pascalschen Dreieck in der n -ten Zeile an der k -ten Stelle von links, wobei die Zeilen- und Stellenzahl jeweils bei 0 beginnen.

Wir sehen, dass im Pascalschen Dreieck die Summe zweier nebeneinanderstehender Zahlen gleich der Zahl direkt unter diesen Zahlen ist. In Formeln:

$$\binom{n}{k} = \binom{n-1}{k-1} + \binom{n-1}{k}. \quad (1.9)$$

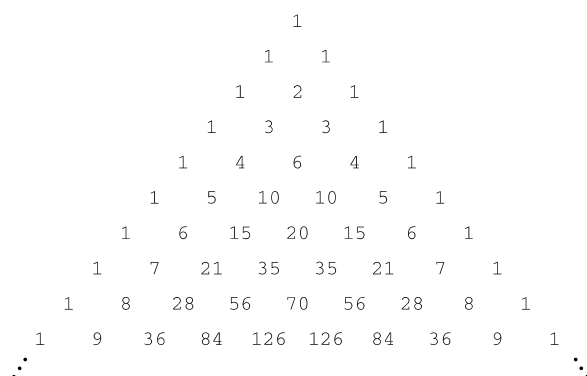


Abbildung 1.1: Das Pascalsche Dreieck

Beweis dazu:

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

□

Der Binomialkoeffizient hat tatsächlich noch eine weitere Bedeutung:

Satz 1.5.2 (kombinatorische Bedeutung des Binomialkoeffizienten).

Die Anzahl der k -elementigen Teilmengen einer n -elementigen Menge $\{a_1, \dots, a_n\}$ ist gleich $\binom{n}{k}$.

Beweis: Es sei C_k^n die Anzahl der k -elementigen Mengen von $\{a_1, \dots, a_n\}$. Wir beweisen den Satz durch vollständige Induktion über die Anzahl n der Elemente.

$n = 1$: $C_0^1 = C_1^1 = \binom{1}{0} = \binom{1}{1} = 1$, da $\{a_1\}$ nur eine nullelementige Teilmenge $\emptyset$ und die einelementige Teilmenge $\{a_1\}$ besitzt.

$n \rightarrow n+1$: Es sei $C_k^n = \binom{n}{k}$ schon bewiesen. Da $C_0^{n+1} = 1 = \binom{n+1}{0}$ und $C_{n+1}^{n+1} = 1 = \binom{n+1}{n+1}$, genügt es, den Fall $1 \leq k \leq n$ zu behandeln.

Die k -elementigen Teilmengen von $\{a_1, \dots, a_{n+1}\}$ zerfallen in zwei Klassen K_0 und K_1 , wobei K_0 alle Teilmengen umfasse, die a_{n+1} nicht enthalten, und K_1 alle Teilmengen, die a_{n+1} enthalten.

Es gehören also genau die k -elementigen Teilmengen von $\{a_1, \dots, a_n\}$ zu K_0 . Derer gibt es nach Induktionsvoraussetzung $\binom{n}{k}$.

Eine Teilmenge gehört genau dann zu K_1 , wenn man sie als Vereinigung von $\{a_{n+1}\}$ mit einer $(k-1)$ -elementigen Teilmenge von $\{a_1, \dots, a_n\}$ darstellen kann. Es gibt also insbesondere genauso viele Teilmengen, die zu K_1 gehören, wie $(k-1)$ -elementige Teilmengen von $\{a_1, \dots, a_n\}$, also nach Induktionsvoraussetzung genau $\binom{n}{k-1}$. Wir haben also

$$C_k^{n+1} = \underbrace{\binom{n}{k}}_{|K_0|} + \underbrace{\binom{n}{k-1}}_{|K_1|} = \binom{n+1}{k}.$$

Damit ist der Schritt von n auf $n+1$ gezeigt, und die Behauptung des Satzes folgt. $\square$

Beispiel 1.5.3 (Kombinationen beim Lotto „6 aus 49“).

Die Anzahl der sechselementigen Teilmengen aus $\{1, \dots, 49\}$ ist

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

Die Chance, im Lotto 6 Richtige zu haben, ist also ungefähr 1 : 14 Millionen.

Kapitel 2

Lineare Algebra I

In der Linearen Algebra geht es um Räume, Vektoren, Matrizen. Sie ist Grundlage für fast alle Gebiete der angewandten Mathematik. Der wesentliche Grund dafür ist die Tatsache, dass sich viele Phänomene mit sogenannten **Linearen Modellen** gut beschreiben lassen, die ein wichtiger Gegenstand der Linearen Algebra sind.

Beispiel 2.0.4 (Bleiaufnahme im Körper).

Frage: Wieviel Blei lagert sich in Blut und Knochen ein (nach Batschelet et al., J. Math. Biology, Vol 8, pp. 15-23, 1979)? Wir sammeln einige Tatsachen über Blei im Körper, und basteln daraus danach ein einfaches lineares Modell.

- Man nimmt jeden Tag ca. $50 \mu\text{g}$ Blei über Lungen und Haut auf, die ins Blut gehen.
- 0,4 % des Bleis im Blut werden jeden Tag in die Knochen eingelagert.
- 2 % des Bleis im Blut werden jeden Tag wieder ausgeschieden.
- 0,004 % des Bleis in den Knochen gehen jeden Tag wieder ins Blut zurück.

Wenn b_j die Bleimenge im Blut am j ten Tag ist, und k_j die in den Knochen, dann können wir die Bleientwicklung von Tag zu Tag durch die folgenden zwei Gleichungen beschreiben:

$$\begin{array}{rcll} k_{j+1} & = & k_j & + 4 \cdot 10^{-3} b_j - 4 \cdot 10^{-5} k_j \\ b_{j+1} & = & b_j & + \underbrace{50 \mu\text{g}}_{\text{Aufnahme}} - \underbrace{4 \cdot 10^{-3} b_j}_{\text{vom Blut in die Knochen}} - \underbrace{2 \cdot 10^{-2} b_j}_{\text{Ausscheidung}} + \underbrace{4 \cdot 10^{-5} k_j}_{\text{von den Knochen ins Blut}} \end{array}$$

Dieses Modell erlaubt uns, zu simulieren, wie sich die Bleikonzentration in Blut und Knochen in einem Individuum in Zukunft verhalten wird. Wir können uns aber z.B. auch fragen, ob es einen Gleichgewichtszustand mit $b_{j+1} = b_j$ und $k_{j+1} = k_j$ gibt, ob dieser sich von selbst einstellt, wenn ja, wie schnell er sich einstellt etc.

Auf all diese Fragen geben Methoden aus der Linearen Algebra eine Antwort. Die Suche nach einem Gleichgewichtswert ist z.B. äquivalent zum Finden zweier Unbekannter b und k , für die gilt:

$$\begin{array}{rcll} 0 & = & & + 4 \cdot 10^{-3} b - 4 \cdot 10^{-5} k \\ 0 & = & + 50 \mu\text{g} & - 4 \cdot 10^{-3} b - 2 \cdot 10^{-2} b + 4 \cdot 10^{-5} k \end{array}$$

Dies ist ein einfaches Beispiel für ein **lineares Gleichungssystem**. In der Praxis tauchen solche Systeme nicht nur mit zwei Unbekannten, sondern leicht mit Hunderten oder Tausenden von Unbekannten auf, und es hilft, wenn man gelernt hat, die Übersicht zu behalten, und in der Lage ist, sie schnell mit Hilfe eines Computers zu lösen.

2.1 Mengen und Abbildungen

2.1.1 Mengen

- **Mengen** sind Zusammenfassungen von wohlunterschiedenen Elementen zu einem Ganzen. Beispiele $\mathbb{N} = \{0, 1, 2, \dots\}$, $\mathbb{Z} = \{\dots, -1, 0, 1, 2, \dots\}$.
- Die **leere Menge** $\{\}$ wird auch mit dem Symbol $\emptyset$ bezeichnet.
- Wir sagen „ A ist **Teilmenge** von B “, falls jedes Element von A auch Element von B ist und schreiben in diesem Fall: $A \subset B$. Es gilt für jede Menge A , dass $\emptyset \subset A$ und $A \subset A$.
- Die **Schnittmenge** von A und B ist die Menge der Elemente, die sowohl in A als auch in B enthalten sind und wird mit $A \cap B$ („ A geschnitten mit B “) bezeichnet.
- Die **Vereinigungsmenge** von A und B ist die Menge aller Elemente, die in A oder in B (oder in beiden Mengen) enthalten sind und wird mit $A \cup B$ („ A vereinigt mit B “) bezeichnet.
- Die **Differenzmenge** $A \setminus B$ („ A ohne B “) ist die Menge aller Elemente aus A , die nicht in B sind. Beispiel: $\mathbb{N} \setminus \{0\} = \{1, 2, \dots\}$.

2.1.2 Das kartesische Produkt

Was ist ein Paar von zwei Elementen? Es besteht aus einem ersten Element a und einem zweiten Element b , und wir bezeichnen das Paar mit (a, b) . Zwei Paare sind nur dann gleich, wenn sowohl das erste als auch das zweite Element übereinstimmen. Es gilt z.B. $(3, 4) \neq (4, 3)$. Wir definieren uns nun die Menge aller Paare aus zwei Mengen A und B .

Definition 2.1.1 (Kartesisches Produkt zweier Mengen).

Sind A und B Mengen, so heißt die Menge $A \times B$ („ A kreuz B “)

$$A \times B := \{(a, b) \mid a \in A, b \in B\}$$

das **kartesische Produkt** der beiden Mengen, das in Abbildung 2.1 illustriert ist.

Ein Beispiel ist z.B. die Menge $\mathbb{R} \times \mathbb{R}$, die man auch $\mathbb{R}^2$ nennt. Man kann auch das kartesische Produkt aus mehr als zwei Mengen bilden.

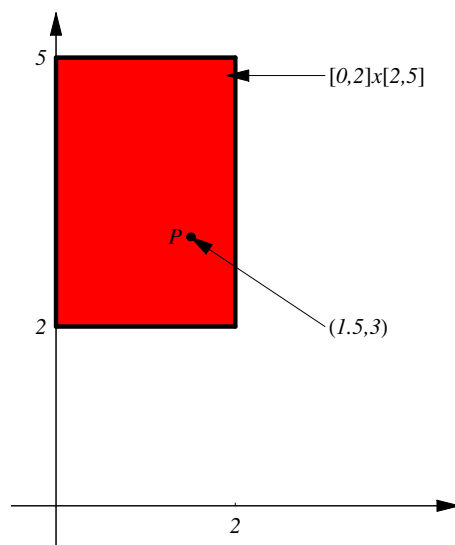


Abbildung 2.1: Das kartesische Mengenprodukt $[0, 2] \times [2, 5]$ und das Paar $(1.5, 3) \in [0, 2] \times [2, 5]$.

Definition 2.1.2 (n -Tupel und kartesisches Mengenprodukt).

Seien $A_1, A_2, \dots, A_n$ Mengen, und $a_1 \in A_1, \dots, a_n \in A_n$. Wir nennen die geordnete Zusammenfassung $(a_1, a_2, \dots, a_n)$ ein **n -Tupel**. Das **kartesische Produkt** der Mengen ist durch

$$A_1 \times A_2 \times \dots \times A_n := \{(a_1, a_2, \dots, a_n) \mid a_1 \in A_1, a_2 \in A_2, \dots, a_n \in A_n\}$$

definiert.

Achtung: n -Tupel sind nur dann gleich, wenn sie zum einen gleich viele Komponenten haben, und zum anderen in jeder Komponente übereinstimmen. Es gilt aber z.B. $(1, 0) \neq (1, 0, 0)$ und $(1, 0, 0) \neq (0, 1, 0)$.

Ein wichtiges Beispiel ist die Menge $\mathbb{R}^n = \underbrace{\mathbb{R} \times \dots \times \mathbb{R}}_{n\text{-mal}}$ aller n -Tupel von reellen Zahlen.

2.1.3 Abbildungen

Definition 2.1.3 (Abbildung, Funktion).

Sind X, Y Mengen, so heißt eine Vorschrift f , die jedem $x \in X$ ein $y \in Y$ zuordnet, eine **Abbildung** oder **Funktion** von X nach Y . Das einem x zugeordnete Element y nennt man $f(x)$. Man schreibt:

$$\begin{aligned} f : X &\rightarrow Y \\ x &\mapsto f(x) \end{aligned}$$

Definition 2.1.4 (Graph einer Abbildung).

Die Menge $\{(x, y) \in X \times Y \mid y = f(x)\}$ heißt der **Graph** von f .

Definition 2.1.5 (Bild, Urbild, Einschränkung einer Abbildung).

Seien $M \subset X$ und $N \subset Y$. Dann heißt

$$f(M) := \{y \in Y \mid \exists x \in M : y = f(x)\}$$

das **Bild von M** , und

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

das **Urbild von N** .

Desweiteren ist $f|_M : M \rightarrow Y \quad x \mapsto f(x)$ die **Einschränkung von f auf M** (vergleiche Abbildung 2.2).

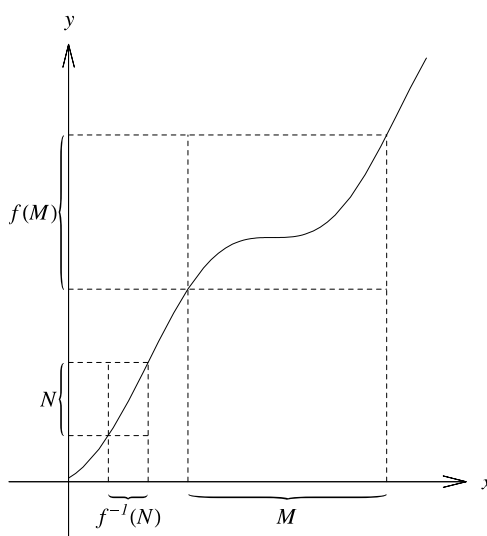


Abbildung 2.2: Bild $f(M)$ der Menge M unter der Abbildung f , und Urbild $f^{-1}(N)$ der Menge N .

Wichtig sind auch die folgenden Begriffe: eine Abbildung $f : X \rightarrow Y$ heißt

- **surjektiv** : $\Leftrightarrow \forall y \in Y \exists x \in X : y = f(x)$. „Für alle y in Y gibt es (mindestens) ein Element x in X , für das gilt: $y = f(x)$ “
- **injektiv** : $\Leftrightarrow \forall x, x' \in X : f(x) = f(x') \Rightarrow x = x'$. „Immer wenn zwei Elemente aus X auf den gleichen Wert abgebildet werden, sind sie gleich.“
- **bijektiv**, wenn f zugleich surjektiv und injektiv ist. Man kann zeigen, dass dies gleichbedeutend ist mit „Jedes Element aus Y ist Bild von genau einem Element aus X “.

Wir sammeln noch ein paar Eigenschaften von Abbildungen.

- Man kann zwei Abbildungen $f_1 : X_1 \rightarrow Y_1$ und $f_2 : X_2 \rightarrow Y_2$ hintereinanderausführen, wenn die Mengen Y_1 und X_2 gleich sind: Man schreibt dann

$$\begin{aligned} f_2 \circ f_1 : X_1 &\longrightarrow Y_2, \\ x &\longmapsto (f_2 \circ f_1)(x) := f_2(f_1(x)), \end{aligned}$$

und man bezeichnet $f_2 \circ f_1$ als die **Verknüpfung** oder **Verkettung** oder auch **Komposition** der zwei Abbildungen.

Achtung: bei Berechnung von $(f_2 \circ f_1)(x)$ wird zuerst f_1 und dann f_2 ausgeführt.

- Die so genannte **Identität auf A** ist eine Abbildung, die jedem Element einer Menge A genau das selbe Element zuordnet:

$$\begin{aligned} \text{Id}_A : A &\longrightarrow A \\ a &\longmapsto a. \end{aligned}$$

Die Identität auf A ist bijektiv.

- Für jede bijektive Abbildung $f : A \rightarrow B$ gibt es eine **Umkehrabbildung** $f^{-1} : B \rightarrow A$ mit den Eigenschaften $f \circ f^{-1} = \text{Id}_B$ und $f^{-1} \circ f = \text{Id}_A$. Achtung: die Umkehrabbildung gibt es *nur* für bijektive Abbildungen, sonst ist sie nicht definiert!

2.2 Reelle Vektorräume

2.2.1 Der $\mathbb{R}^n$ als reeller Vektorraum

Mit Zahlen aus $\mathbb{R}$ kann man rechnen, man kann sie addieren, multiplizieren etc. Was kann man mit n -Tupeln reeller Zahlen $(x_1, x_2, \dots, x_n)$ machen? Wir fassen sie in Zukunft selbst wieder als Variable auf, die wir auch **Vektor** nennen, z.B. $x = (x_1, x_2, \dots, x_n)$ oder $y = (y_1, y_2, \dots, y_n)$. Wir können nun die Addition $x + y$ zweier gleich langer n -Tupel $x \in \mathbb{R}^n$ und $y \in \mathbb{R}^n$ definieren. (Im Folgenden ist n einfach eine feste natürliche Zahl).

Definition 2.2.1 (Vektoraddition).

$$(x_1, \dots, x_n) + (y_1, \dots, y_n) := (x_1 + y_1, \dots, x_n + y_n).$$

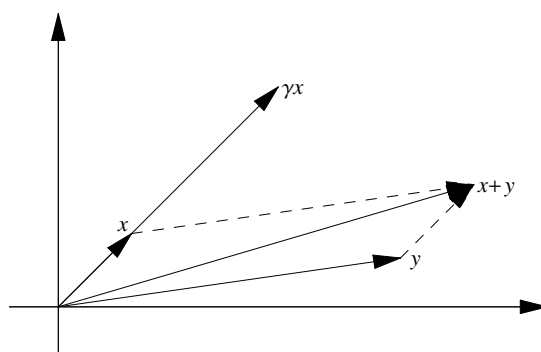


Abbildung 2.3: Summe $x + y$ von zwei Vektoren im $\mathbb{R}^2$ und die Streckung γx von x um den Faktor γ .

Man beachte, dass die Vektoraddition zwar das gleiche Symbol „+“ wie die normale Addition reeller Zahlen benutzt, aber etwas davon Verschiedenes ist, nämlich eine Abbildung

$$\begin{aligned} + : \mathbb{R}^n \times \mathbb{R}^n &\longrightarrow \mathbb{R}^n, \\ (x, y) &\longmapsto x + y. \end{aligned}$$

Des Weiteren definieren wir eine Multiplikation eines Vektors $x \in \mathbb{R}^n$ mit einem Skalar $\lambda \in \mathbb{R}$.

Definition 2.2.2 (Skalarmultiplikation).

$$\lambda (x_1, \dots, x_n) := (\lambda x_1, \dots, \lambda x_n).$$

Die Skalarmultiplikation ist eine Abbildung

$$\begin{aligned} \cdot : \mathbb{R} \times \mathbb{R}^n &\longrightarrow \mathbb{R}^n, \\ (\lambda, x) &\longmapsto \lambda x. \end{aligned}$$

Vektoraddition und Skalarmultiplikation sind in Abbildung 2.3 illustriert. Unter Beachtung der Rechenregeln für reelle Zahlen ergibt sich:

1. Für $x, y, z \in \mathbb{R}^n$ gilt

$$(x + y) + z = x + (y + z) \quad \text{[Assoziativgesetz]}.$$

2. $\forall x, y \in \mathbb{R}^n$ gilt

$$x + y = y + x \quad \text{[Kommutativgesetz]}.$$

3. Der Nullvektor $0 := (0, \dots, 0)$ ist das **neutrale Element** der Vektoraddition:

$$v + 0 = v \quad \forall v \in \mathbb{R}^n.$$

4. Sei für $v = (v_1, \dots, v_n)$ das **Negative** durch $-v := (-v_1, \dots, -v_n)$ definiert. Dann gilt

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

5. $\forall x, y \in \mathbb{R}^n$ und $\lambda, \mu \in \mathbb{R}$ gilt

$$\begin{aligned} (\lambda\mu)x &= \lambda(\mu x), \\ 1x &= x, \\ \lambda(x + y) &= \lambda x + \lambda y, \\ (\lambda + \mu)x &= \lambda x + \mu x. \end{aligned}$$

Wir beweisen als Übung nur die letzte Gleichung:

$$\begin{aligned} (\lambda + \mu)x &= ((\lambda + \mu)x_1, \dots, (\lambda + \mu)x_n) \\ &= (\lambda x_1 + \mu x_1, \dots, \lambda x_n + \mu x_n) \\ &= (\lambda x_1, \dots, \lambda x_n) + (\mu x_1, \dots, \mu x_n) \\ &= \lambda x + \mu x. \end{aligned}$$

2.2.2 Allgemeine Vektorräume

Wir haben nun die Menge $\mathbb{R}^n$ mit zwei Rechenoperationen, der Vektoraddition und der Skalarmultiplikation, ausgestattet. Dies erlaubt uns, mit den n -Tupeln reeller Zahlen auf eine bestimmte Weise zu rechnen, die auch in vielen anderen Bereichen der Mathematik nützlich ist. Deshalb verallgemeinern Mathematiker die soeben beobachteten Rechenregeln, und sagen: Jede Menge V , mit deren Elementen man eine Addition und eine Skalarmultiplikation durchführen kann, nennen wir einen **reellen Vektorraum**.

Definition 2.2.3 (Reeller Vektorraum).

Ein Tripel $(V, +, \cdot)$, bestehend aus einer Menge V , einer Abbildung

$$\begin{aligned} + : V \times V &\longrightarrow V, \\ (x, y) &\longmapsto x + y, \end{aligned}$$

und einer Abbildung

$$\begin{aligned} \cdot : \mathbb{R} \times V &\longrightarrow V, \\ (\lambda, x) &\longmapsto \lambda x, \end{aligned}$$

heißt **reeller Vektorraum**, wenn die folgenden acht **Vektorraumaxiome** gelten:

1. $\forall x, y, z \in V : \quad (x + y) + z = x + (y + z).$
2. $\forall x, y \in V : \quad x + y = y + x.$
3. $\exists 0 \in V \forall x \in V : \quad 0 + x = x.$
4. $\forall x \in V \exists y \in V : \quad x + y = 0.$
5. $\forall x \in V, \lambda, \mu \in \mathbb{R} : \quad (\lambda\mu)x = \lambda(\mu x).$
6. $\forall x \in V : \quad 1x = x.$
7. $\forall x, y \in V, \lambda \in \mathbb{R} : \quad \lambda(x + y) = \lambda x + \lambda y.$
8. $\forall x \in V, \lambda, \mu \in \mathbb{R} : \quad (\lambda + \mu)x = \lambda x + \mu x.$

2.2.3 Untervektorräume

Manche Teilmengen eines Vektorraums bilden selbst wieder einen Vektorraum. Solche Teilmengen heißen **Untervektorräume**.

Definition 2.2.4 (Untervektorraum).

Sei $(V, +, \cdot)$ ein reeller Vektorraum und $W \subset V$ eine Teilmenge. W heißt **Untervektorraum** von V , falls die folgenden **Untervektorraumaxiome** gelten:

UV1: $W \neq \emptyset$

UV2: $\forall v, w \in W : v + w \in W$, d.h. W ist gegenüber der Addition abgeschlossen.

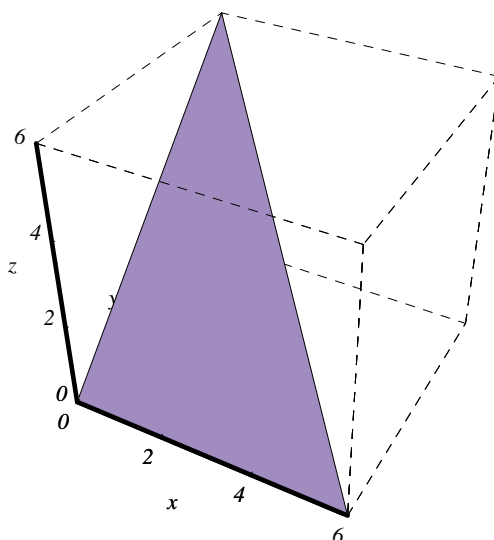


Abbildung 2.4: Einen zweidimensionalen Untervektorraum im $\mathbb{R}^3$ kann man sich als gekippte Ebene vorstellen.

UV3: $\forall v \in W, \lambda \in \mathbb{R} : \lambda v \in W$, d.h. W ist gegenüber der Skalarmultiplikation abgeschlossen.

In Abbildung 2.4 ist ein zweidimensionaler Untervektorraum im $\mathbb{R}^3$ skizziert.

Lemma 2.2.5 (Jeder Untervektorraum ist ein Vektorraum).

Ist V ein reeller Vektorraum und $W \subset V$ ein Untervektorraum, so ist W mit der aus V induzierten Addition und Skalarmultiplikation selbst wieder ein reeller Vektorraum

Beweis: Kommutativ- und Assoziativgesetz gelten natürlich, da sie in V gelten. Der Nullvektor 0 liegt in W , da wegen (UV1) ein $v \in V$ existiert und somit wegen (UV3) gilt, dass $0 = 0v \in W$. Zu jedem $v \in V$ ist wegen (UV3) auch $-v = (-1)v \in V$. Das inverse Element liegt also auch in W . Damit ist W ein Vektorraum. $\square$

2.3 *Gruppen, Körper, Vektorräume

In diesem Abschnitt wollen wir noch einige Konzepte einführen, die zwar grundlegend für die Mathematik sind, aber an dieser Stelle nicht unbedingt nötig für das Verständnis der Linearen Algebra sind. Wem die axiomatische Formulierung des Vektorraums bereits genug der Abstraktion ist, der kann diesen Abschnitt getrost überspringen; wem diese Art des Verallgemeinerns gefällt, der bekommt hier mehr davon.

2.3.1 Gruppen

Der Begriff der Gruppe findet sich in allen möglichen Bereichen der Mathematik wieder, da er sehr allgemein ist. Man kann an Hand nur sehr weniger Voraussetzungen schon viele Dinge

beweisen, und es ist ein ganzer Zweig der Mathematik, die **Gruppentheorie**, aus der folgenden Definition entsprungen.

Definition 2.3.1 (Gruppe).

1. Eine **Gruppe** ist ein Paar $(G, \cdot)$, bestehend aus einer Menge G und einer Verknüpfung „ $\cdot$ “:

$$\begin{aligned} \cdot : G \times G &\rightarrow G \\ (a, b) &\mapsto a \cdot b, \end{aligned}$$

mit folgenden Eigenschaften (**Gruppenaxiomen**):

G1: (**Assoziativgesetz**)

$$\forall a, b, c \in G \quad (a \cdot b) \cdot c = a \cdot (b \cdot c). \quad (2.1)$$

G2: Es existiert ein **neutrales Element**:

$$\exists e \in G \forall a \in G \quad e \cdot a = a \cdot e = a. \quad (2.2)$$

G3: Zu jedem Element existiert ein **inverses Element**:

$$\forall a \in G \exists b \in G \quad a \cdot b = b \cdot a = e. \quad (2.3)$$

2. Gilt für eine Gruppe $(G, \cdot)$ zusätzlich noch das **Kommutativgesetz**,

$$\forall a, b \in G \quad a \cdot b = b \cdot a, \quad (2.4)$$

so wird sie **kommutative** oder auch **abelsche Gruppe** genannt.

Bemerkung 2.3.2 (Notation der Verknüpfung).

Man lässt in der Notation das Verknüpfungszeichen „ $\cdot$ “ häufig weg, schreibt also z.B. ab anstatt $a \cdot b$, so wie bei der gewöhnlichen Multiplikation. In anderen Fällen, gerade bei kommutativen Gruppen, benutzt man aber gerne auch ein anderes Verknüpfungszeichen, nämlich „ $+$ “. Warum, wird am besten anhand einiger Beispiele deutlich.

Beispiele für Gruppen

- Die Menge $\mathbb{R}$ der reellen Zahlen bildet zusammen mit der üblichen Addition eine kommutative Gruppe. Das neutrale Element ist die Zahl 0.
- Die Menge $\mathbb{R} \setminus \{0\}$ der reellen Zahlen ohne die Null bildet zusammen mit der üblichen Multiplikation eine kommutative Gruppe. Das neutrale Element ist die Zahl 1.
- Die Menge $\mathbb{Z} = \{\dots, -1, 0, 1, 2, \dots\}$ bildet zusammen mit der üblichen Addition eine kommutative Gruppe, mit neutralem Element 0. Warum ist $\mathbb{Z}$ mit der Multiplikation keine Gruppe? Warum ist die Menge $\mathbb{N} = \{0, 1, 2, \dots\}$ weder mit der Addition noch mit der Multiplikation eine Gruppe?
- Ein ganz anderes Beispiel ist die Menge $\text{Bij}(A)$ aller bijektiven Abbildungen $f : A \rightarrow A$ einer nichtleeren Menge A auf sich selbst, zusammen mit der Abbildungs-Verknüpfung, denn wenn f und g in $\text{Bij}(A)$ sind, so ist auch $f \circ g$ wieder in $\text{Bij}(A)$. Das neutrale Element dieser Gruppe ist die Identität Id_A , das Inverse zu f ist gerade die Umkehrabbildung f^{-1} .

2.3.2 Körper

Das zweite Konzept verallgemeinert das Konzept der reellen Zahlen, mit denen man wie gewohnt rechnen kann, zu dem Begriff des Körpers.

Definition 2.3.3 (Körper).

Ein **Körper** ist ein Tripel $(\mathbb{K}, +, \cdot)$, bestehend aus einer Menge $\mathbb{K}$ und zwei Verknüpfungen $+$ und $\cdot$ auf $\mathbb{K}$, d.h. einer Abbildung (Addition)

$$\begin{aligned} + : \mathbb{K} \times \mathbb{K} &\longrightarrow \mathbb{K}, \\ (a, b) &\longmapsto a + b, \end{aligned}$$

und einer Abbildung (Multiplikation)

$$\begin{aligned} \cdot : \mathbb{K} \times \mathbb{K} &\longrightarrow \mathbb{K}, \\ (a, b) &\longmapsto a \cdot b, \end{aligned}$$

mit den Eigenschaften (**Körperaxiomen**):

K1: $(\mathbb{K}, +)$ ist eine kommutative Gruppe

Das neutrale Element ist wie mit 0 bezeichnet.

K2: $(\mathbb{K} \setminus \{0\}, \cdot)$ ist eine kommutative Gruppe Das neutrale Element ist wie mit 1 bezeichnet.

K3: $a \cdot (b + c) = (a \cdot b) + (a \cdot c) \quad \forall a, b, c \in \mathbb{K}$ [**Distributivgesetz**].

Beispiele für Körper

- Die Menge der reellen Zahlen $\mathbb{R}$ mit Addition und Multiplikation bildet einen Körper.
- Die Menge der rationalen Zahlen $\mathbb{Q}$ mit Addition und Multiplikation bildet einen Körper.
- Wir werden in Kapitel 5 die Menge $\mathbb{C}$ der komplexen Zahlen kennenlernen, die mit einer Addition und Multiplikation ausgestattet ist und auch einen Körper bildet.

2.3.3 Allgemeine Vektorräume

Die Definition des Begriffs des Körpers erlaubt uns nun, noch einen allgemeineren Typ von Vektorraum zu definieren. Es werden einfach die reellen Zahlen in der Definition des reellen Vektorraums durch die Elemente irgendeines Körpers ersetzt. Außerdem können wir mit Hilfe des Gruppenbegriffs die ersten Axiome kürzer schreiben.

Definition 2.3.4 ($\mathbb{K}$ -Vektorraum).

Sei $\mathbb{K}$ ein Körper. Ein **$\mathbb{K}$ -Vektorraum** ist ein Tripel $(\mathbb{V}, +, \cdot)$ bestehend aus einer Menge $\mathbb{V}$,

einer Verknüpfung „+“ mit

$$\begin{aligned} + : \mathbb{V} \times \mathbb{V} &\rightarrow \mathbb{V} \\ (v, w) &\mapsto v + w, \end{aligned}$$

einer Verknüpfung „·“ mit

$$\begin{aligned} \cdot : \mathbb{K} \times \mathbb{V} &\rightarrow \mathbb{V}, \\ (\lambda, \mu) &\mapsto \lambda v, \end{aligned}$$

für die die folgenden **Vektorraumaxiome** gelten:

V1: $(\mathbb{V}, +)$ ist eine abelsche Gruppe [Das neutrale Element 0 heißt **Nullvektor**, das zu einem $v \in \mathbb{V}$ inverse Element heißt der zu v **negative Vektor**].

V2: $\forall v, w \in \mathbb{V}, \lambda, \mu \in \mathbb{K}$ gilt:

- (a) $(\lambda\mu)v = \lambda(\mu v)$,
- (b) $1v = v$,
- (c) $\lambda(v + w) = (\lambda v) + (\lambda w)$,
- (d) $(\lambda + \mu)v = (\lambda v) + (\mu v)$.

Statt $\mathbb{K}$ -Vektorraum sagt man auch **Vektorraum über $\mathbb{K}$** . Wir haben schon gesehen, dass die n -Tupel reeller Zahlen einen reellen Vektorraum, also einen Vektorraum über $\mathbb{R}$ bilden.

Beispiel 2.3.5 (Vektorraum von Abbildungen).

Sei X eine Menge, $\mathbb{K}$ ein Körper, etwa $X = \mathbb{R}$ und $\mathbb{K} = \mathbb{R}$. Sei $F(X, \mathbb{K})$ die Menge aller Abbildungen von X nach $\mathbb{K}$. Ein $f \in F(\mathbb{R}, \mathbb{R})$ ist etwa $f(x) = x^2$.

Durch die Addition

$$\begin{aligned} (f, g) &\mapsto f + g && \text{für } f, g \in F(X, \mathbb{K}), \text{ mit} \\ (f + g)(x) &:= f(x) + g(x), \end{aligned}$$

und die Skalarmultiplikation

$$\begin{aligned} (\lambda, f) &\mapsto \lambda f, \\ (\lambda f)(x) &:= \lambda(f(x)), \end{aligned}$$

wird $(F(X, \mathbb{K}), +, \cdot)$ zu einem $\mathbb{K}$ -Vektorraum.

Das Inverse von $f \in F$ ist durch

$$(-f)(x) := -f(x)$$

definiert.

2.4 Skalarprodukt, euklidische Norm und Vektorprodukt

Wir führen nun für den recht anschaulichen Vektorraum $\mathbb{R}^3$ einige geometrische Begriffe ein. Unser Ziel ist es u.a., eine Distanz zwischen zwei Elementen (Vektoren) des $\mathbb{R}^3$ festzulegen. Bis auf das Vektorprodukt lassen sich alle Begriffe auf naheliegende Weise auf den $\mathbb{R}^n$ verallgemeinern. Im Kapitel 6 gehen wir darauf ausführlich ein.

Definition 2.4.1 (Standard-Skalarprodukt in $\mathbb{R}^3$).

Seien $x, y \in \mathbb{R}^3$. Der Wert

$$\langle x, y \rangle := x_1y_1 + x_2y_2 + x_3y_3$$

heißt das **Standard-Skalarprodukt** von x und y . Dadurch ist eine Abbildung von $\mathbb{R}^3 \times \mathbb{R}^3$ nach $\mathbb{R}$ definiert.

Für $x, y, z \in \mathbb{R}^3, \lambda \in \mathbb{R}$ gilt:

1. $\langle x + y, z \rangle = \langle x, z \rangle + \langle y, z \rangle$.
2. $\langle \lambda x, y \rangle = \lambda \langle x, y \rangle$.
3. $\langle x, y \rangle = \langle y, x \rangle$.
4. $\langle x, x \rangle \geq 0$ und $\langle x, x \rangle = 0 \Leftrightarrow x = 0$.

2.4.1 Norm und Distanz

Mit Hilfe des Skalarproduktes können wir nun einige Begriffe definieren, die sich anschaulich interpretieren lassen.

Definition 2.4.2 (Euklidische Norm eines Vektors).

Sei $x \in \mathbb{R}^3$. Dann heißt

$$\|x\| := \sqrt{\langle x, x \rangle} = \sqrt{x_1^2 + x_2^2 + x_3^2}$$

die **euklidische Norm** oder auch die **euklidische Länge** von x .

Es gilt: $\|x\| = 0 \Leftrightarrow x = 0$, und $\|\lambda x\| = |\lambda| \cdot \|x\|$. Jedem Vektor wird durch die Norm ein Skalar zugeordnet. Anschaulich gilt: Je größer die Norm von x , desto weiter ist x vom Ursprung entfernt. Die Norm ermöglicht es uns nun auch, einen Abstand zwischen Vektoren zu definieren.

Definition 2.4.3 (Distanz von Vektoren).

Für $x, y \in \mathbb{R}^3$ ist $\|x - y\|$ die **Distanz** oder auch der **Abstand** zwischen x und y .

Es gilt für alle $x, y, z \in \mathbb{R}^3$:

1. $\|x - y\| \geq 0$ und $(\|x - y\| = 0 \Leftrightarrow x = y)$.

2. $\|x - y\| = \|y - x\|$.
3. $\|x - z\| \leq \|x - y\| + \|y - z\|$. (Dreiecksungleichung)

Nur der letzte Punkt, die Dreiecksungleichung, ist nicht offensichtlich und bedarf eines Beweises, den wir am Ende des folgenden Abschnitts geben.

2.4.2 Eigenschaften des Skalarproduktes

Seien $x, y, z, \in \mathbb{R}^3$. Dann gelten folgende Gleichungen und Ungleichungen:

1. Verallgemeinerter Satz des Pythagoras:

$$\|x + y\|^2 = \|x\|^2 + \|y\|^2 + 2\langle x, y \rangle.$$

Falls x, y orthogonal zueinander sind (s. Definition 2.4.4), dann gilt sogar $\|x + y\|^2 = \|x\|^2 + \|y\|^2$.

Beweis: Wir verwenden die nach Definition 2.4.1 aufgelisteten Rechenregeln des Skalarprodukts.

$$\begin{aligned} \|x + y\|^2 &= \langle x + y, x + y \rangle \\ &= \langle x, x \rangle + \langle x, y \rangle + \langle y, x \rangle + \langle y, y \rangle \\ &= \|x\|^2 + 2\langle x, y \rangle + \|y\|^2. \end{aligned}$$

2. Cauchy-Schwarzsche Ungleichung:

$$|\langle x, y \rangle| \leq \|x\| \cdot \|y\|.$$

Beweis: Ist $y = 0$, so sind linke und rechte Seite gleich 0, d.h. die Behauptung stimmt. Es genügt, $y \neq 0$ zu betrachten.

Sei $\lambda := \langle y, y \rangle, \mu := -\langle x, y \rangle$ Dann ist

$$\begin{aligned} 0 &\leq \langle \lambda x + \mu y, \lambda x + \mu y \rangle \\ &= \lambda^2 \langle x, x \rangle + 2\lambda\mu \langle x, y \rangle + \mu^2 \langle y, y \rangle \\ &= \lambda(\langle x, x \rangle \langle y, y \rangle - 2\langle x, y \rangle^2 + \langle x, y \rangle^2) \\ &= \lambda(\langle x, x \rangle \langle y, y \rangle - \langle x, y \rangle^2) \end{aligned}$$

wegen $\lambda > 0$ folgt daraus

$$\langle x, y \rangle^2 \leq \langle x, x \rangle \langle y, y \rangle$$

und wegen der Monotonie der Quadratwurzel die Behauptung. $\square$

In Kapitel 6 geben wir einen geometrischen Beweis der Cauchy-Schwarz-Ungleichung, s. Korollar 6.2.1.

3. Dreiecksungleichung:

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

Beweis:

$$\begin{aligned} \|x + y\|^2 &= \|x\|^2 + 2\langle x, y \rangle + \|y\|^2 \\ &\leq \|x\|^2 + 2\|x\| \cdot \|y\| + \|y\|^2 \\ &= (\|x\| + \|y\|)^2. \end{aligned}$$

Dabei haben wir im vorletzten Schritt die Cauchy-Schwarzsche Ungleichung verwendet. Also ist $\|x + y\|^2 \leq (\|x\| + \|y\|)^2$ und wegen der Monotonie der Wurzel $\|x + y\| \leq \|x\| + \|y\|$. $\square$

Aus der Dreiecksungleichung für die Norm folgt direkt auch die Dreiecksungleichung für die Distanz von Vektoren aus Definition 2.4.3, indem man x und y durch $x - y$ und $y - z$ ersetzt.

4. Man kann das Skalarprodukt $\langle x, y \rangle$ anschaulich interpretieren, wenn man sich die beiden Vektoren in der von ihnen aufgespannten Ebene ansieht. Mit dem Winkel ϕ zwischen ihnen in dieser Ebene gilt nämlich (ohne Beweis, illustriert in Abbildung 2.5):

$$\langle x, y \rangle = \cos(\phi) \|x\| \|y\|.$$

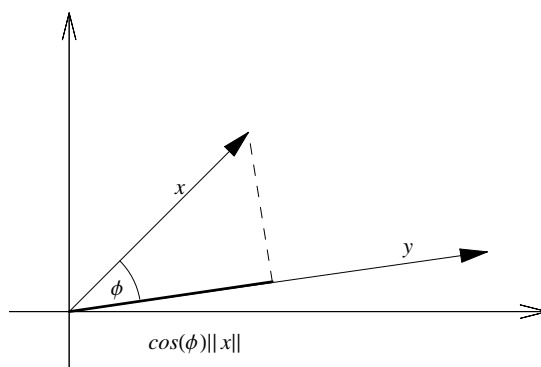


Abbildung 2.5: Das Skalarprodukt der Vektoren x und y graphisch veranschaulicht.

Die letzte Interpretation des Skalarprodukts motiviert folgende Definition:

Definition 2.4.4 (Orthogonalität).

Zwei Vektoren $x, y \in \mathbb{R}^3$ heißen **orthogonal** bzw. **senkrecht zueinander**, wenn

$$\langle x, y \rangle = 0.$$

2.4.3 Das Vektorprodukt

Für die Physik ist ein weiteres Produkt zwischen Vektoren wichtig, das allerdings nur im $\mathbb{R}^3$, also dem physikalischen Raum, definiert ist: das sogenannte Vektorprodukt.

Definition 2.4.5 (Vektorprodukt).

Für $x, y \in \mathbb{R}^3$ ist

$$x \times y := \begin{pmatrix} x_2 y_3 - x_3 y_2 \\ x_3 y_1 - x_1 y_3 \\ x_1 y_2 - x_2 y_1 \end{pmatrix}$$

das **Vektorprodukt** von x und y .

Das Vektorprodukt hat für alle $x, y \in \mathbb{R}^3$ folgende Eigenschaften:

1. $\langle x, x \times y \rangle = 0$ und $\langle y, x \times y \rangle = 0$, d.h. $x \times y$ ist senkrecht zu x und y .
2. Wenn ϕ der (positive) Winkel zwischen x und y ist, dann gilt

$$\|x \times y\| = \sin(\phi) \|x\| \|y\|.$$

Dies kann man so interpretieren, dass $\|x \times y\|$ der Flächeninhalt des durch x und y aufgespannten Parallelogramms ist.

2.5 Lineare Unabhängigkeit, Basis und Dimension

In diesem Abschnitt wollen wir versuchen, ein Maß für die „Größe“ eines Vektorraumes zu finden. Das geeignete Maß hierfür ist die **Dimension** eines Vektorraumes, deren Definition wir uns jetzt Schritt für Schritt nähern wollen. Zunächst definieren wir uns einige in diesem Zusammenhang wichtige Begriffe.

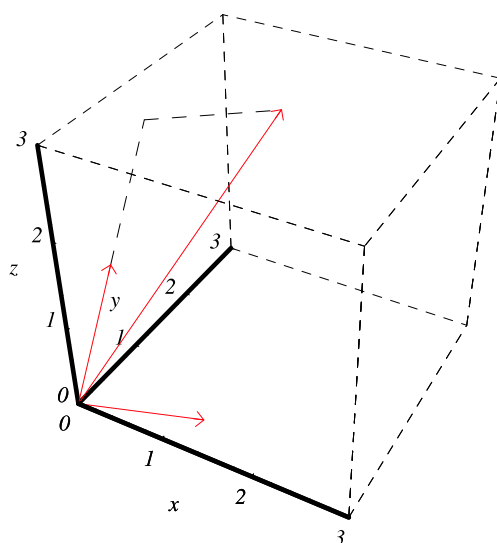
Definition 2.5.1 (Linearkombination).

Sei $(V, +, \cdot)$ ein reeller Vektorraum, und seien $(v_1, \dots, v_r)$, $r \geq 1$ Vektoren aus V . Ein $x \in V$ heißt **Linearkombination** aus $(v_1, \dots, v_r)$, falls es $\lambda_1, \dots, \lambda_r \in \mathbb{R}$ gibt, so dass

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

Man sagt auch: „ x lässt sich aus $v_1, \dots, v_r$ linear kombinieren.“

Mit Hilfe des Begriffs der Linearkombination lässt sich nun folgende Menge definieren:

Abbildung 2.6: Linearkombination im $\mathbb{R}^3$ **Definition 2.5.2** (Spann, lineare Hülle).

Der **Spann der Vektoren** $v_1, \dots, v_r$,

$$\text{Spann}(v_1, \dots, v_r) := \{\lambda_1 v_1 + \dots + \lambda_r v_r \mid \lambda_1, \dots, \lambda_r \in \mathbb{R}\},$$

ist die Menge aller Vektoren aus V , die sich aus $v_1, \dots, v_r$ linear kombinieren lassen. $\text{Spann}(v_1, \dots, v_r)$ heißt auch **der durch** $v_1, \dots, v_r$ **aufgespannte Raum** oder die **lineare Hülle** der Vektoren $v_1, \dots, v_r$. Man kann leicht zeigen, dass $\text{Spann}(v_1, \dots, v_r)$ selbst wieder ein Vektorraum ist.

Intuitiv liegt es nahe, die Dimension mit Hilfe des Spanns zu definieren. Man kann z.B. zwei Vektoren verwenden, um den $\mathbb{R}^2$ aufzuspannen, denn

$$\mathbb{R}^2 = \text{Spann} \left(\begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right).$$

Wir werden sehen, dass die Anzahl der zum Aufspannen eines Raumes benötigten Vektoren tatsächlich die Dimension des Raumes festlegt. Ein Problem ist allerdings, dass man auch mehr Vektoren als nötig nehmen könnte, z.B.

$$\mathbb{R}^2 = \text{Spann} \left(\begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right).$$

Einer der Vektoren, z.B. der dritte, ist überflüssig, da er selbst wieder als Linearkombination der anderen dargestellt werden kann. Um solche Fälle ausschließen zu können, definieren wir uns die folgenden beiden Begriffe.

Definition 2.5.3 (Lineare Abhängigkeit).

Ein r -Tupel von Vektoren $(v_1, \dots, v_r)$ heißt linear abhängig, wenn mindestens einer der Vektoren als Linearkombination der anderen dargestellt werden kann.

Wichtig für unsere Zwecke ist nun aber gerade der Fall, dass die Vektoren *nicht* linear abhängig sind. Es lässt sich zeigen, dass die Verneinung der linearen Abhängigkeit gerade durch die folgende Definition gegeben ist:

Definition 2.5.4 (Lineare Unabhängigkeit).

Sei V ein reeller Vektorraum. Die Vektoren $v_1, \dots, v_r \in V$ heißen **linear unabhängig** (siehe Abbildung 2.7), falls gilt:

Sind $\lambda_1, \dots, \lambda_r \in \mathbb{R}$ und ist

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

so folgt notwendig

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

Man sagt auch: „Der Nullvektor lässt sich nur trivial aus den Vektoren $v_1, \dots, v_r$ linear kombinieren.“ Mit Hilfe des Begriffs der linearen Unabhängigkeit lässt sich nun erst der Begriff der

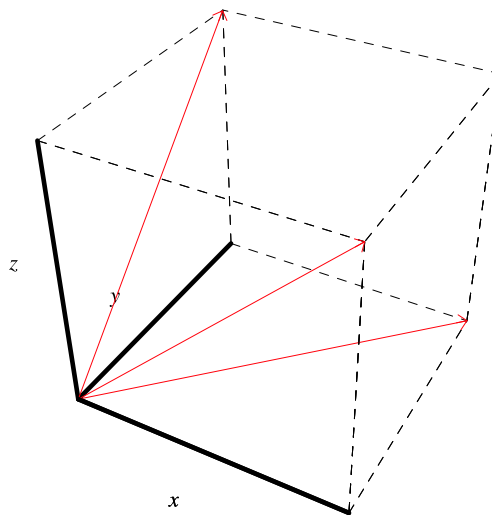


Abbildung 2.7: Drei linear unabhängige Vektoren

Basis, und damit endlich auch die Dimension eines Vektorraumes definieren.

Definition 2.5.5 (Basis).

Die Vektoren $v_1, \dots, v_r$ aus einem reellen Vektorraum V bilden eine **Basis** von V , falls gilt:

B1: $\text{Spann}(v_1, \dots, v_r) = V$,

B2: Die Vektoren $v_1, \dots, v_r$ sind linear unabhängig.

Definition 2.5.6 (Dimension).

Hat ein Vektorraum V eine endliche Basis $(v_1, \dots, v_r)$ mit r Elementen, so definiert man seine **Dimension** als

$$\dim V := r.$$

Diese Definition der Dimension eines Vektorraums mit Hilfe irgendeiner beliebigen Basis ist auf Grund des folgenden Satzes gerechtfertigt.

Satz 2.5.7. Je zwei endliche Basen eines reellen Vektorraumes haben die gleiche Anzahl von Elementen.

Beispiel 2.5.8 (Eine Basis des $\mathbb{R}^n$).

Sei $e_i := (0, \dots, 0, 1, 0, \dots, 0)$, $1 \leq i \leq n$, wobei die „1“ an der i -ten Stelle steht.

Sind $\lambda_1, \dots, \lambda_n \in \mathbb{R}$ Skalare mit $\lambda_1 e_1 + \dots + \lambda_n e_n = 0$, so folgt wegen $\lambda_1 e_1 + \dots + \lambda_n e_n = (\lambda_1, \dots, \lambda_n)$, dass $\lambda_1 = \dots = \lambda_n = 0$ sein muß. Also sind $e_1, \dots, e_n$ linear unabhängig und B2 ist somit erfüllt.

Sei $v \in V = \mathbb{R}^n$ ein beliebiger Vektor, mit $v = (v_1, \dots, v_n)$. Wegen $v = v_1 e_1 + \dots + v_n e_n$ ist auch B1 erfüllt, daher bilden die n Vektoren $(e_1, \dots, e_n)$ eine Basis des $\mathbb{R}^n$, die sogenannte **kanonische Basis**.

2.5.1 Basis-Isomorphismen

Mit Hilfe einer Basis kann jeder n -dimensionale Vektorraum mit dem $\mathbb{R}^n$ identifiziert werden: Sei V ein beliebiger Vektorraum und $\mathcal{B} = (v_1, \dots, v_n)$, $v_i \in V$ eine Basis von V . Dann gibt es genau eine bijektive Abbildung

$$\begin{aligned} \phi_{\mathcal{B}} : \mathbb{R}^n &\rightarrow V, \\ (x_1, \dots, x_n) &\mapsto \phi_{\mathcal{B}}(x) := x_1 v_1 + \dots + x_n v_n. \end{aligned}$$

Die Abbildung $\phi_{\mathcal{B}}$ nennt man auch **Basis-Isomorphismus** oder **Koordinationsystem** und $x = (x_1, \dots, x_n) \in \mathbb{R}^n$ den **Koordinatenvektor** von $v = x_1 v_1 + \dots + x_n v_n \in V$ bezüglich $\mathcal{B}$. Es gilt $v = \phi_{\mathcal{B}}(x)$ und $x = \phi_{\mathcal{B}}^{-1}(v)$. Die Abbildung $\phi_{\mathcal{B}}$ hat neben der Bijektivität eine weitere wichtige Eigenschaft, sie ist *linear*. Mit linearen Abbildungen werden wir uns im folgenden sehr intensiv beschäftigen.

2.6 Lineare Abbildungen

Definition 2.6.1 (Lineare Abbildung, Vektorraumhomomorphismus).

Seien V und W zwei reelle Vektorräume, und $F : V \rightarrow W$ eine Abbildung. F heißt **linear**, falls $\forall v, w \in V, \lambda \in \mathbb{R}$ gilt:

$$\text{L1: } F(v + w) = F(v) + F(w),$$

$$\text{L2: } F(\lambda v) = \lambda F(v).$$

Eine lineare Abbildung wird auch **Homomorphismus** genannt. Die Menge aller linearen Abbildungen von V nach W wird mit $\text{Hom}(V, W)$ bezeichnet.

Wir können die Eigenschaften (L1) und (L2) auch zusammenfassen zu

$$\forall v, w \in V, \lambda, \mu \in \mathbb{R} : \quad F(\lambda v + \mu w) = \lambda F(v) + \mu F(w),$$

und in Worten interpretieren als „ F ist mit den auf V und vorgegebenen Verknüpfungen $+$ und $\cdot$ verträglich.“ Die folgenden Eigenschaften einer linearen Abbildung F sind leicht zu zeigen:

1. $F(0) = 0$ und $F(v - w) = F(v) - F(w) \quad \forall v, w \in V$.
2. Sind $v_1, \dots, v_r$ Vektoren in V , so gilt:
 - (a) Sind $(v_1, \dots, v_r)$ linear abhängig in V , so sind $(F(v_1), \dots, F(v_r))$ linear abhängig in W .
 - (b) Sind $(F(v_1), \dots, F(v_r))$ linear unabhängig in W , so sind $(v_1, \dots, v_r)$ linear unabhängig in V .
3. Sind $V' \subset V$ und $W' \subset W$ Untervektorräume, so sind auch $F(V') \subset W$ und $F^{-1}(W') \subset V$ Untervektorräume.
4. $\dim F(V) \leq \dim V$.

Beweis:

1.

$$\begin{aligned} \text{Es gilt } F(0) &= F(0 + 0) \\ &\stackrel{(L1)}{=} F(0) + F(0). \end{aligned}$$

Subtraktion von $F(0)$ auf beiden Seiten liefert

$$F(0) = 0$$

Die zweite Gleichung folgt aus

$$\begin{aligned} F(v - w) &= F(v + (-w)) \\ &\stackrel{(L1)}{=} F(v) + F(-w) \\ &\stackrel{(L2)}{=} F(v) - F(w). \end{aligned}$$

2. (a) Gibt es $i_1, \dots, i_k \in \{1, \dots, r\}$ und $\lambda_1, \dots, \lambda_k \in \mathbb{R} \setminus \{0\}$ mit $\lambda_1 v_{i_1} + \dots + \lambda_k v_{i_k} = 0$, so ist auch

$$\lambda_1 F(v_{i_1}) + \dots + \lambda_k F(v_{i_k}) = 0.$$

- (b) Wegen der Äquivalenz von $A \Rightarrow B$ mit $\neg B \Rightarrow \neg A$ ist diese Aussage äquivalent zu 2.(a).

3. Wir beweisen nur $F(V') \subset W$. Wegen $0 \in V'$ ist $0 = F(0) \in F(V')$. Sind $w, w' \in F(V')$, so gibt es $v, v' \in V'$ mit $F(v) = w$ und $F(v') = w'$. Also ist $w + w' = F(v) + F(v') = F(v + v') \in F(V')$, denn $v + v' \in V'$.

Ist andererseits $\lambda \in \mathbb{R}$ und $w \in F(V')$, so ist $\lambda w = \lambda F(v) = F(\lambda v) \in F(V')$, denn $\lambda v \in V'$. Also ist $F(V')$ ist Untervektorraum von W . Der Beweis $F^{-1}(W') \subset V$ geht analog (freiwillige Übung).

4. folgt aus 2. □

Beispiele für lineare Abbildungen

- Basis-Isomorphismen wie in Abschnitt 2.5.1 sind lineare Abbildungen. Allgemein nennt man übrigens jede bijektive lineare Abbildung **Isomorphismus**.
- Die **Nullabbildung** $0 : V \rightarrow \{0\}$ und die Identität auf V sind linear. Achtung: Für ein $0 \neq v_0 \in W$ ist die konstante Abbildung $F : V \rightarrow W, F(v) = v_0 \quad \forall v \in V$ *nicht* linear.
- Das wichtigste Beispiel ist sicher die folgende Form einer linearen Abbildung. Seien für $1 \leq i \leq m$ und $1 \leq j \leq n$ reelle Zahlen a_{ij} gegeben, und sei $F : \mathbb{R}^n \rightarrow \mathbb{R}^m$ durch

$$F(x_1, \dots, x_n) := \left(\sum_{j=1}^n a_{1j} x_j, \quad \dots, \quad \sum_{j=1}^n a_{mj} x_j \right)$$

gegeben. Durch *einfaches Einsetzen* kann gezeigt werden, dass F linear ist. Tatsächlich hat jede lineare Abbildung von $\mathbb{R}^n \rightarrow \mathbb{R}^m$ diese Gestalt.

Eine Verallgemeinerung des letzten Beispiels ist fundamental für das Verständnis linearer Abbildungen und das Arbeiten mit ihnen.

Satz 2.6.2 (Matrixdarstellung einer Linearen Abbildung).

Seien V und W Vektorräume mit Basen $\mathcal{A} = (v_1, \dots, v_n)$ und $\mathcal{B} = (w_1, \dots, w_m)$, und seien für $1 \leq i \leq m$ und $1 \leq j \leq n$ die reellen Zahlen a_{ij} gegeben. Dann ist durch

$$\begin{aligned} F(v_1) &:= a_{11}w_1 + \dots + a_{m1}w_m \\ &\vdots \\ F(v_n) &:= a_{1n}w_1 + \dots + a_{mn}w_m \end{aligned} \quad (2.5)$$

eine lineare Abbildung $F : V \rightarrow W$ eindeutig definiert. Umgekehrt lassen sich zu **jeder** linearen Abbildung F eindeutig bestimmte Zahlen a_{ij} ($1 \leq i \leq m$ und $1 \leq j \leq n$) finden, die (2.5) erfüllen.

Das heißt, bei gegebenen Basen der Räume V und W kann jede lineare Abbildung $F : V \rightarrow W$ durch eine *Zahlentabelle* eindeutig repräsentiert werden. Diese Zahlentabelle nennt man auch die **darstellende Matrix** der Abbildung F zu den Basen $\mathcal{A}$ und $\mathcal{B}$, und bezeichnet sie manchmal mit dem Symbol $M_{\mathcal{B}}^{\mathcal{A}}(F)$.

Beweis: Zunächst zeigen wir, dass F durch die Gleichungen (2.5) wohldefiniert ist: Sei $v \in V$, so gibt es eindeutig bestimmte und $\lambda_1, \dots, \lambda_n \in \mathbb{R}$, so dass

$$v = \lambda_1 v_1 + \dots + \lambda_n v_n.$$

Da F linear ist, gilt

$$F(v) = \lambda_1 F(v_1) + \dots + \lambda_n F(v_n),$$

und die Vektoren $F(v_1), \dots, F(v_n)$ sind durch (2.5) eindeutig definiert.

Wir beweisen nun die Umkehrung, dass sich zu jeder linearen Abbildung F eine darstellende Matrix finden läßt. Da sich jeder Vektor $w \in W$ eindeutig als Linearkombination aus $(w_1, \dots, w_m)$ darstellen läßt, gilt auch für die Bilder der Basisvektoren $F(v_j) \in W$, dass es für $j = 1, \dots, n$ eindeutig bestimmte Skalare $a_{1j}, \dots, a_{mj}$ gibt, so dass

$$F(v_j) = a_{1j}w_1 + \dots + a_{mj}w_m.$$

□

2.6.1 Bild, Rang und Kern

Definition 2.6.3 (Rang).

Ist $F : V \rightarrow W$ eine lineare Abbildung so bezeichnen wir mit

$$\begin{aligned} \text{Bild}(F) &:= F(V) = \{F(v) \mid v \in V\} && \text{das \textbf{Bild} von } F \\ \text{Rang}(F) &:= \dim \text{Bild}(F) && \text{den \textbf{Rang} von } F, \text{ und mit} \\ \text{Ker}(F) &:= F^{-1}(0) = \{v \in V \mid F(v) = 0\} && \text{den \textbf{Kern} von } F. \end{aligned}$$

Die Mengen $\text{Bild}(F)$ und $\text{Ker}(F)$ sind selbst wieder Vektorräume, und es gilt der folgende Satz (ohne Beweis):

Satz 2.6.4 (Dimensionsformel).

$$\dim(V) = \dim \text{Bild}(F) + \dim \text{Ker}(F).$$

Für Bild und Kern gelten folgende Eigenschaften:

- $\text{Rang}(F) \leq \dim V$
- $\text{Ker}(F) = \{0\} \Leftrightarrow F$ ist injektiv,
- $\text{Rang}(F) = \dim W \Leftrightarrow F$ ist surjektiv,
- $\dim V = \dim W$ und $\text{Ker}(F) = \{0\} \Leftrightarrow F$ ist bijektiv.

2.7 Matrizen

Das Arbeiten mit linearen Abbildungen wird wesentlich vereinfacht durch die Verwendung von Matrizen. Wir führen hier zunächst einfach die Matrizen und ihre Rechenregeln ein, und kommen dann im nächsten Abschnitt auf ihre Bedeutung in der linearen Algebra zu sprechen.

Definition 2.7.1 (Matrix).

Eine Tabelle reeller Zahlen mit m Zeilen und n Spalten nennen wir eine **reelle** $(m \times n)$ -**Matrix**. Man schreibt

$$A = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix}$$

mit Koeffizienten $a_{ij} \in \mathbb{R}$ für $1 \leq i \leq m$ und $1 \leq j \leq n$.

Die Menge aller reellen $(m \times n)$ -Matrizen bezeichnet man mit $\mathbb{R}^{m \times n}$ („ $\mathbb{R}$ hoch m kreuz n “).

Definition 2.7.2 (Addition und Skalarmultiplikation).

Wir können auf der Menge $\mathbb{R}^{m \times n}$ eine Addition und Skalarmultiplikation einführen, ebenso wie

wir es für Vektoren getan hatten:

$$\begin{aligned}
 A + B &= \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix} + \begin{pmatrix} b_{11} & \cdots & b_{1n} \\ \vdots & & \vdots \\ b_{m1} & \cdots & b_{mn} \end{pmatrix} \\
 &:= \begin{pmatrix} a_{11} + b_{11} & \cdots & a_{1n} + b_{1n} \\ \vdots & & \vdots \\ a_{m1} + b_{m1} & \cdots & a_{mn} + b_{mn} \end{pmatrix}, \\
 \lambda A &= \lambda \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix} \\
 &:= \begin{pmatrix} \lambda a_{11} & \cdots & \lambda a_{1n} \\ \vdots & & \vdots \\ \lambda a_{m1} & \cdots & \lambda a_{mn} \end{pmatrix}.
 \end{aligned}$$

Definition 2.7.3 (Transponierte Matrix).

Ist $A = (a_{ij}) \in \mathbb{R}^{m \times n}$ so sei $a_{ji}^T := a_{ij} \in \mathbb{R}^{n \times m}$ und die Matrix $A^T := (a_{ji}^T) \in \mathbb{R}^{n \times m}$ (lies „ A transponiert“) heißt die **zu A transponierte Matrix**.

Beispiel 2.7.4.

$$\begin{pmatrix} 6 & 2 & 3 \\ 9 & 0 & 4 \end{pmatrix}^T = \begin{pmatrix} 6 & 9 \\ 2 & 0 \\ 3 & 4 \end{pmatrix}.$$

Definition 2.7.5 (Matrizenmultiplikation).

Ist $A = (a_{ij}) \in \mathbb{R}^{m \times n}$ und $B = (b_{ij}) \in \mathbb{R}^{n \times r}$ so ist das **Produkt von A und B** , $A \cdot B = (c_{ik})$, durch

$$c_{ik} := \sum_{j=1}^n a_{ij} b_{jk} = a_{i1} b_{1k} + a_{i2} b_{2k} + \cdots + a_{in} b_{nk}$$

für $i = 1, \dots, m$ und $k = 1, \dots, r$ definiert. Es gilt $A \cdot B \in \mathbb{R}^{m \times r}$, also ist die Multiplikation als Abbildung

$$\begin{aligned}
 \mathbb{R}^{m \times n} \times \mathbb{R}^{n \times r} &\rightarrow \mathbb{R}^{m \times r}, \\
 (A, B) &\mapsto A \cdot B,
 \end{aligned}$$

aufzufassen.

Achtung: Die Spaltenzahl n von A muß mit der Zeilenzahl von B übereinstimmen. $A \cdot B$ hat so viele Zeilen wie A und so viele Spalten wie B :

$$\begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{i1} & \cdots & a_{in} \\ \vdots & & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix} \cdot \begin{pmatrix} b_{11} & \cdots & b_{1k} & \cdots & b_{1r} \\ \vdots & & \vdots & & \vdots \\ b_{n1} & \cdots & b_{nk} & \cdots & b_{nr} \end{pmatrix} = \begin{pmatrix} c_{11} & \cdots & \cdots & \cdots & c_{1r} \\ \vdots & & \vdots & & \vdots \\ \vdots & \cdots & c_{ik} & \cdots & \vdots \\ \vdots & & \vdots & & \vdots \\ c_{m1} & \cdots & \cdots & \cdots & a_{mr} \end{pmatrix}.$$

So entsteht c_{ik} aus der i -ten Zeile von A und der k -ten Spalte von B .

Beispiel 2.7.6.

$$\begin{pmatrix} 6 & 2 & 3 \\ 9 & 0 & 4 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 & 2 & 2 \\ 2 & 4 & 1 & 0 \\ 3 & 5 & 0 & 0 \end{pmatrix} = \begin{pmatrix} 19 & 23 & 14 & 12 \\ 21 & 20 & 18 & 18 \end{pmatrix}.$$

2.7.1 Rechenregeln für Matrizen

- Für $A, B \in \mathbb{R}^{m \times n}$ und $\lambda \in \mathbb{R}$ gilt (Beweis durch Einsetzen):

$$\begin{aligned} (A + B)^T &= A^T + B^T, \\ (\lambda A)^T &= \lambda A^T, \\ (A^T)^T &= A, \\ (AB)^T &= B^T A^T. \end{aligned}$$

- Man beachte: Für die Matrixmultiplikation gilt im allgemeinen $AB \neq BA$. Es ist etwa

$$\begin{aligned} \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} &= \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}, \\ \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} &= \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}. \end{aligned}$$

- Eine spezielle Matrix ist die n -reihige **Einheitsmatrix**

$$\mathbb{I}_n := \begin{pmatrix} 1 & & 0 \\ & \ddots & \\ 0 & & 1 \end{pmatrix} \in \mathbb{R}^{n \times n}.$$

Es gilt

$$\forall A \in \mathbb{R}^{n \times m} : \quad A \mathbb{I}_m = \mathbb{I}_n A = A. \quad (2.6)$$

- Für die Matrizen $A, A' \in \mathbb{R}^{m \times n}$, $B, B' \in \mathbb{R}^{n \times r}$ und $\lambda \in \mathbb{R}$ gilt:

$$\begin{aligned} A(B + B') &= AB + AB', \\ (A + A')B &= AB + A'B \quad [\text{Distributivgesetz}], \\ A(\lambda B) &= (\lambda A)B = \lambda(AB), \\ (AB)C &= A(BC) \quad [\text{Assoziativgesetz}]. \end{aligned}$$

2.7.2 Von der Matrix zur linearen Abbildung

Wir werden nun sehen, dass die Matrizen einen ganz direkten Zusammenhang mit linearen Abbildungen haben. Alles wird einfacher, wenn wir die Elemente des $\mathbb{R}^n$ jetzt als **Spaltenvektoren** schreiben, also als $(n \times 1)$ -Matrix. Wir schreiben z.B.

$$x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \in \mathbb{R}^n.$$

Dies erlaubt uns, auch die Matrix-Vektor-Multiplikation mit Hilfe der normalen Matrizenmultiplikation auszudrücken, z.B. für eine $(m \times n)$ -Matrix A und $x \in \mathbb{R}^n$ können wir $Ax \in \mathbb{R}^m$ berechnen als

$$\begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} a_{11}x_1 + \dots + a_{1n}x_n \\ \vdots \\ a_{m1}x_1 + \dots + a_{mn}x_n \end{pmatrix}$$

Mit dieser Konvention können wir den Zusammenhang zwischen Matrizen und linearen Abbildungen in sehr kompakter Form ausdrücken.

Satz 2.7.7 (Matrix einer linearen Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^m$).

Sei A eine reelle $(m \times n)$ -Matrix. Dann ist durch

$$\begin{aligned} F : \mathbb{R}^n &\rightarrow \mathbb{R}^m, \\ x &\mapsto F(x) := Ax, \end{aligned}$$

eine lineare Abbildung F definiert. Umgekehrt gibt es zu jeder linearen Abbildung $F : \mathbb{R}^n \rightarrow \mathbb{R}^m$ eine Matrix $A \in \mathbb{R}^{m \times n}$, so dass $\forall x \in \mathbb{R}^n : F(x) = Ax$.

Wegen

$$F(e_j) = Ae_j = \begin{pmatrix} a_{11} & \dots & a_{1j} & \dots & a_{1n} \\ \vdots & & \vdots & & \vdots \\ a_{m1} & \dots & a_{mj} & \dots & a_{mn} \end{pmatrix} \begin{pmatrix} 0 \\ \vdots \\ 1 \\ \vdots \\ 0 \end{pmatrix} = \begin{pmatrix} a_{1j} \\ \vdots \\ a_{mj} \end{pmatrix}$$

gilt:

Die Spaltenvektoren von A sind die Bilder der kanonischen Basisvektoren.

Beispiel 2.7.8. Sei $F : \mathbb{R}^3 \rightarrow \mathbb{R}^2$ durch $F(x_1, x_2, x_3) = (3x_1 + 2x_3, x_2 + 2x_3)$ gegeben. Dann wird F dargestellt durch $A = \begin{pmatrix} 3 & 0 & 2 \\ 0 & 1 & 2 \end{pmatrix}$.

Mit diesem Zusammenhang zwischen linearen Abbildungen und Matrizen können wir nun auch Begriffe wie Bild, Rang und Kern einer Abbildung direkt auf Matrizen übertragen. Es gilt für eine Matrix $A = (a_1, a_2, \dots, a_n) \in \mathbb{R}^{m \times n}$ mit Spaltenvektoren $a_1, \dots, a_n$:

- $\text{Bild}(A) := \{Ax \in \mathbb{R}^m \mid x \in \mathbb{R}^n\} = \text{Spann}(a_1, \dots, a_n)$
- $\text{Rang}(A) := \dim \text{Bild}(A)$, die maximale Anzahl linear unabhängiger Spaltenvektoren.
- $\text{Ker}(A) := \{x \in \mathbb{R}^n \mid Ax = 0\}$.

Wegen der Dimensionsformel (Satz 2.6.4) gilt: $\dim \text{Ker}(A) = n - \text{Rang}(A)$.

Man kann durch Nachrechnen auch den folgenden sehr wichtigen Satz zeigen, der im Nachhinein die Definition der Matrixmultiplikation rechtfertigt:

Satz 2.7.9 (Matrixprodukt als Verknüpfung linearer Abbildungen). Ist $A = (a_{ij}) \in \mathbb{R}^{m \times n}$ und $B = (b_{ij}) \in \mathbb{R}^{n \times r}$ und $a : \mathbb{R}^n \rightarrow \mathbb{R}^m$ und $b : \mathbb{R}^r \rightarrow \mathbb{R}^n$ die durch A und B dargestellten linearen Abbildungen. Dann gilt für ihre Verknüpfung $a \circ b$:

$$(a \circ b)(x) = ABx.$$

Die Matrixmultiplikation beschreibt die Verknüpfung zweier linearer Abbildungen.

2.7.3 Inversion von Matrizen

Definition 2.7.10 (Regularität und Singularität einer quadratischen Matrix).

Eine (quadratische) Matrix $A \in \mathbb{R}^{n \times n}$ heißt **invertierbar** oder auch **regulär**, falls es eine Matrix $A^{-1} \in \mathbb{R}^{n \times n}$ gibt mit:

$$AA^{-1} = A^{-1}A = \mathbb{I}_n.$$

Falls A nicht regulär ist, dann heißt A **singulär**.

Satz 2.7.11 (Bedingungen für Regularität einer quadratischen Matrix).

Sei $F : \mathbb{R}^n \rightarrow \mathbb{R}^m$ eine lineare Abbildung und sei A die darstellende Matrix von F , d.h. $F(x) = Ax$. Dann sind folgende Aussagen einander äquivalent:

- F ist ein Isomorphismus (also bijektiv).
- $n = m = \text{Rang}(F)$.

(c) Die darstellende Matrix A ist regulär.

In diesem Falle gilt:

$$F^{-1}(y) = A^{-1}y \quad \forall y \in \mathbb{R}^m.$$

Eine bijektive lineare Abbildung F bezeichnet man als **Isomorphismus**.

Die Umkehrabbildung wird durch die inverse Matrix dargestellt.

Es gibt noch eine wichtige Rechenregel für inverse Matrizen:

Satz 2.7.12. Seien $A, B \in \mathbb{R}^{n \times n}$ zwei invertierbare Matrizen. Dann ist auch ihr Matrixprodukt AB invertierbar, und es gilt

$$(AB)^{-1} = B^{-1}A^{-1}.$$

Ein Algorithmus zum Invertieren

Wir werden nun einen Algorithmus zur Berechnung der Inversen einer regulären Matrix kennenlernen.

Definition 2.7.13 (Elementare Zeilenumformungen).

U_1 : Multiplikation der i -ten Zeile mit $\lambda \neq 0$.

U_2 : Addition des λ -fachen der j -ten Zeile zur i -ten Zeile.

U_3 : Vertauschen der i -ten und der j -ten Zeile.

Satz 2.7.14. Elementare Umformungen U_1, U_2 und U_3 ändern den Rang einer Matrix $A \in \mathbb{R}^{n \times n}$ nicht.

Beispiel 2.7.15 (Für elementare Zeilenumformung).

Die Matrizen

$$\begin{pmatrix} 3 & 7 & 3 \\ 6 & 2 & 0 \\ 9 & 1 & 1 \end{pmatrix} \rightarrow \begin{pmatrix} 3 & 7 & 3 \\ 9 & 1 & 1 \\ 6 & 2 & 0 \end{pmatrix} \rightarrow \begin{pmatrix} 3 & 7 & 3 \\ 12 & 8 & 4 \\ 6 & 2 & 0 \end{pmatrix} \rightarrow \begin{pmatrix} 9 & 9 & 3 \\ 12 & 8 & 4 \\ 6 & 2 & 0 \end{pmatrix}$$

haben den gleichen Rang. Es wurden erst die Zeilen 2 und 3 vertauscht, dann zur neuen Zeile 2 Zeile 1 addiert, dann zur Zeile 1 Zeile 3 addiert.

Satz 2.7.16 (Berechnung der inversen Matrix).

Man kann eine reguläre Matrix S durch elementare Umformungen in die Einheitsmatrix überführen. Wenn man „parallel dazu“ die gleichen Umformungen auf die Einheitsmatrix anwendet, erhält man aus der umgeformten Einheitsmatrix die Inverse von S .

Beispiel 2.7.17 (Für die Berechnung der Inversen).

$$\begin{aligned}
 S &= \begin{pmatrix} 3 & -2 \\ -1 & 1 \end{pmatrix}, \\
 \left(\begin{array}{cc|cc} 3 & -2 & 1 & 0 \\ -1 & 1 & 0 & 1 \end{array} \right) &\rightarrow \left(\begin{array}{cc|cc} 1 & 0 & 1 & 2 \\ -1 & 1 & 0 & 1 \end{array} \right) \rightarrow \left(\begin{array}{cc|cc} 1 & 0 & 1 & 2 \\ 0 & 1 & 1 & 3 \end{array} \right), \\
 \Rightarrow S^{-1} &= \begin{pmatrix} 1 & 2 \\ 1 & 3 \end{pmatrix}
 \end{aligned}$$

2.8 Lineare Gleichungssysteme

Ein wichtiges Ziel der linearen Algebra besteht darin, Aussagen über die Lösungen eines linearen Gleichungssystems

$$\begin{array}{ccccccc}
 a_{11}x_1 + & \cdots & + a_{1n}x_n & = & b_1 \\
 \vdots & & \vdots & & \vdots \\
 a_{m1}x_1 + & \cdots & + a_{mn}x_n & = & b_m
 \end{array}$$

mit Koeffizienten a_{ij} und b_i im $\mathbb{R}$ zu machen. Wir können ein solches Gleichungssystem mit Hilfe einer Matrix $A \in \mathbb{R}^{m \times n}$ und eines Vektors $b \in \mathbb{R}^m$ kurz schreiben als

$$\text{„Finde } x \in \mathbb{R}^n \text{ , so dass } Ax = b.\text{“}$$

Wir suchen die **Lösungsmenge**

$$\text{Lös}(A, b) := \{x \in \mathbb{R}^n \mid Ax = b\}.$$

Als erstes wollen wir untersuchen, wie man ein sogenanntes *homogenes Gleichungssystem* löst, d.h. ein solches von der Form $Ax = 0$.

2.8.1 Homogene lineare Gleichungssysteme

Definition 2.8.1 (Homogenes lineares Gleichungssystem).

Seien $a_{ij} \in \mathbb{R}$ für $i = 1, \dots, m$ und $j = 1, \dots, n$. Das Gleichungssystem

$$\begin{array}{ccccccc}
 a_{11}x_1 & + \dots + & a_{1n}x_n & = & 0 \\
 \vdots & & \vdots & & \\
 a_{m1}x_1 & + \dots + & a_{mn}x_n & = & 0
 \end{array} \tag{2.7}$$

wird **homogenes lineares Gleichungssystem** in den **Unbestimmten** $x_1, \dots, x_n$ mit Koeffizienten in $\mathbb{R}$ genannt. Die Matrix

$$A = \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix}$$

heißt **Koeffizientenmatrix**. Mit $x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$ lässt sich (2.7) kurz auch $Ax = 0$ schreiben. Ein (als Spalte) geschriebener Vektor x heißt **Lösung** von (2.7), falls

$$Ax = 0$$

gilt. Unter dem **Lösungsraum** von (2.7) verstehen wir

$$\text{Lös}(A, 0) = \text{Ker}(A) = \{x \in \mathbb{R}^n \mid Ax = 0\}$$

Satz 2.8.2 ($\text{Lös}(A, 0)$ ist ein Untervektorraum).

Ist $A \in \mathbb{R}^{m \times n}$, so ist der Lösungsraum $\text{Lös}(A, 0)$ des zugehörigen homogenen linearen Gleichungssystems ein Untervektorraum des $\mathbb{R}^n$ mit

$$\dim \text{Lös}(A, 0) = \dim \text{Ker}(A) = n - \text{Rang}(A).$$

Beweis: Die Behauptung folgt direkt aus der Dimensionsformel (Satz 2.6.4).

Lösungsverfahren für lineare Gleichungssysteme

Ein Gleichungssystem zu lösen heißt, ein Verfahren anzugeben, nach dem *alle* Lösungen *explizit* zu erhalten sind. Im Falle eines homogenen linearen Gleichungssystems reicht es, eine Basis $(w_1, \dots, w_k)$ des Kerns zu bestimmen, denn dann folgt

$$\text{Ker}(A) = \text{Spann}(w_1, w_2, \dots, w_k).$$

Das Lösungsverfahren basiert auf folgender Beobachtung:

Lemma 2.8.3 (Äquivalente Gleichungssysteme).

Sei $A \in \mathbb{R}^{m \times n}$, $b \in \mathbb{R}^m$ und $S \in \mathbb{R}^{m \times m}$ eine invertierbare Matrix. Dann haben die beiden linearen Gleichungssysteme

$$Ax = b \quad \text{und} \quad (SA)x = Sb$$

die gleichen Lösungsmengen. Insbesondere haben auch

$$Ax = 0 \quad \text{und} \quad (SA)x = 0$$

die gleichen Lösungsmengen.

Beweis: Ist $Ax = b$, so auch $(SA)x = S \cdot (Ax) = Sb$.

Ist umgekehrt $(SA)x = Sb$, so folgt $Ax = S^{-1}((SA)x) = S^{-1}Sb = b$. $\square$

Wir kennen bereits die elementaren Zeilenumformungen. Sie verändern die Lösungsmenge eines Gleichungssystems nicht, denn Sie haben die folgende wichtige Eigenschaft:

Elementare Zeilenumformungen einer Matrix erfolgen durch Multiplikation von links mit einer invertierbaren Matrix.

Denn seien

- A_1 durch Multiplikation der i -ten Zeile mit λ ($\lambda \neq 0$),
- A_2 durch Addition des λ -fachen der j -ten Zeile zur i -ten Zeile,
- A_3 durch Vertauschen der i -ten mit der j -ten Zeile

aus einer Matrix $A \in \mathbb{R}^{m \times n}$ entstanden, dann gilt:

$$\begin{aligned} A_1 &= S_i(\lambda)A, \\ A_2 &= Q_i^j(\lambda)A, \\ A_3 &= P_i^j A, \end{aligned}$$

wobei $S_i(\lambda), Q_i^j(\lambda), P_i^j \in \mathbb{R}^{m \times m}$:

$$S_i(\lambda) = \begin{pmatrix} 1 & & & & \\ & \ddots & & & \\ & & 1 & & 0 \\ & & & \lambda & \\ & & & & 1 \\ & 0 & & & \ddots & \\ & & & & & 1 \end{pmatrix} \leftarrow i\text{-te Zeile,}$$

$\uparrow$
 $i\text{-te Spalte}$

$$Q_i^j(\lambda) = \begin{pmatrix} 1 & & & & 0 \\ & \ddots & & & \\ & & 1 & \lambda & \\ & & & \ddots & \\ 0 & & & & 1 \end{pmatrix} \leftarrow i\text{-te Zeile,}$$

$\uparrow$
 $j\text{-te Spalte}$

$$P_i^j = \begin{pmatrix} 1 & & & & & & 0 \\ & \ddots & & & & & \\ & & 1 & & & & \\ & & & 0 & & 1 & \\ & & & & 1 & & \\ & & & & & \ddots & \\ & & & & & & 1 & 0 \\ & & & & & & & 1 & \ddots \\ 0 & & & & & & & & & 1 \end{pmatrix} \begin{array}{l} \leftarrow i\text{-te Zeile} \\ \\ \\ \leftarrow j\text{-te Zeile} \end{array}$$

$\uparrow$ $i\text{-te Spalte}$ $\uparrow$ $j\text{-te Spalte}$

Diese Matrizen heissen **Elementarmatrizen**, und sie sind alle invertierbar. Es gilt nämlich

- $S_i(\lambda)^{-1} = S_i(\frac{1}{\lambda})$,
- $Q_i^j(\lambda)^{-1} = Q_i^j(-\lambda)$ und
- $(P_i^j)^{-1} = P_i^j$.

Sei $A \in \mathbb{R}^{m \times n}$ und sei $B \in \mathbb{R}^{m \times n}$ aus A durch elementare Zeilenumformungen entstanden. Dann haben

$$Ax = 0 \quad \text{und} \quad Bx = 0$$

die gleichen Lösungsräume. □

Damit können wir Gleichungssysteme vereinfachen! Zunächst bringen wir A durch elementare Zeilenumformungen auf **Zeilenstufenform**

$$B = \begin{pmatrix} b_{1j_1} & \cdots & & \\ 0 & b_{2j_2} & & \\ \vdots & 0 & \ddots & \\ 0 & \cdots & & b_{rj_r} \\ \vdots & & & 0 & \vdots \end{pmatrix},$$

wobei $r = \text{Rang} A$, also auch $r = \text{Rang} A$, $\dim \text{Ker}(A) = n - r = k$. Das Gleichungssystem $Bx = 0$ wird **reduziertes Gleichungssystem** genannt. Es bleibt eine Basis von $\text{Ker}(B) = \text{Ker}(A)$ zu bestimmen. Zur Vereinfachung sei $j_1 = 1, \dots, j_r = r$, was durch Spaltenvertauschun-

gen von B immer erreicht werden kann. Sei also

$$B = \begin{pmatrix} b_{11} & \cdots & & & \\ & 0 & \ddots & & \\ & \vdots & \ddots & b_{rr} & \cdots \\ & 0 & \cdots & 0 & \cdots & 0 \\ & \vdots & & \vdots & & \vdots \end{pmatrix}.$$

Die Unbekannten $x_{r+1}, \dots, x_n$ unterscheiden sich wesentlich von $x_1, \dots, x_r$, denn erstere sind *frei wählbare Parameter*, und $x_1, \dots, x_r$ werden dadurch festgelegt. Sind also $\lambda_1, \dots, \lambda_k \in \mathbb{R}$ beliebig, so gibt es dazu genau ein

$$x = (x_1, \dots, x_r, \lambda_1, \dots, \lambda_k) \in \text{Ker}(B).$$

Die Berechnung von $x_1, \dots, x_r$ zu vorgegebenen $\lambda_1, \dots, \lambda_k$ geschieht *rekursiv rückwärts*. Die r -te Zeile von B ergibt

$$b_{rr}x_r + b_{r,r+1}\lambda_1 + \dots + b_{rn}\lambda_k = 0$$

und wegen $b_{rr} \neq 0$ ergibt sich hieraus x_r . Analog erhält man aus der $(r-1)$ -ten Zeile x_{r-1} und schließlich aus der ersten Zeile x_1 . Insgesamt erhält man eine lineare Abbildung

$$\begin{aligned} G: \mathbb{R}^k &\rightarrow \mathbb{R}^n \\ (\lambda_1, \dots, \lambda_k) &\mapsto (x_1, \dots, x_r, \lambda_1, \dots, \lambda_k). \end{aligned}$$

Diese Abbildung ist injektiv und ihr Bild ist in $\text{Ker}(A)$ enthalten. Wegen $\dim \text{Ker}(A) = k = \text{Rang}(G)$ ist $\text{Bild}(G) = \text{Ker}(A)$. Ist $(e_1, \dots, e_s)$ die kanonische Basis des $\mathbb{R}^k$, so ist $(G(e_1), \dots, G(e_s))$ eine Basis des Kerns $\text{Ker}(B) = \text{Ker}(A)$.

Beispiel 2.8.4 (Lösen eines linearen Gleichungssystems).

$n = 6, m = 4$

$$\begin{array}{rclcl} x_2 & & +2x_4 - x_5 - 4x_6 & = & 0 \\ & x_3 & -x_4 - x_5 + 2x_6 & = & 0 \\ x_2 & & +2x_4 + x_5 - 2x_6 & = & 0 \\ 2x_3 & -2x_4 - 2x_5 + 4x_6 & = & 0 & \end{array}$$

Koeffizientenmatrix A :

$$A = \begin{pmatrix} 0 & 1 & 0 & 2 & -1 & -4 \\ 0 & 0 & 1 & -1 & -1 & 2 \\ 0 & 1 & 0 & 2 & 1 & -2 \\ 0 & 0 & 2 & -2 & -2 & 4 \end{pmatrix}$$

↓ elementare Zeilenumformungen

$$B = \begin{pmatrix} 0 & 1 & 0 & 2 & -1 & -4 \\ 0 & 0 & 1 & -1 & -1 & 2 \\ 0 & 0 & 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

↓ reduziertes Gleichungssystem

$$\begin{array}{rclcl} x_2 & +2x_4 & -x_5 - 4x_6 & = & 0 \\ x_3 & -x_4 & -x_5 + 2x_6 & = & 0 \\ & & x_5 + x_6 & = & 0 \end{array}$$

Es ist

$$r = \text{Rang}(A) = 3, \quad k = \dim \text{Ker} A = 3.$$

Setze

$$x_1 = \lambda_1, x_4 = \lambda_2, x_6 = \lambda_3$$

Es ist

$$\begin{array}{rclcl} x_5 & = & -x_6 & = & -\lambda_3 \\ x_3 & = & x_4 + x_5 - 2x_6 & = & \lambda_2 - \lambda_3 - 2\lambda_3 \\ & & & = & \lambda_2 - 3\lambda_3 \\ x_2 & = & -2x_4 + x_5 + 4x_6 & = & -2\lambda_2 - \lambda_3 + 4\lambda_3 \\ & & & = & -2\lambda_2 + 3\lambda_3 \end{array}$$

Somit ist der Lösungsraum $\text{Ker}(A)$ Bild der injektiven linearen Abbildung

$$\begin{aligned} G: \mathbb{R}^3 &\rightarrow \mathbb{R}^6, \\ (\lambda_1, \lambda_2, \lambda_3) &\mapsto (\lambda_1, -2\lambda_2 + 3\lambda_3, \lambda_2 - 3\lambda_3, \lambda_2, -\lambda_3, \lambda_3). \end{aligned}$$

Insbesondere ist

$$\begin{aligned} G(1, 0, 0) &= (1, 0, 0, 0, 0, 0) = w_1, \\ G(0, 1, 0) &= (0, -2, 1, 1, 0, 0) = w_2, \\ G(0, 0, 1) &= (0, 3, -3, 0, -1, 1) = w_3, \end{aligned}$$

oder allgemein $\text{Ker}(A) = \text{Spann}(w_1, w_2, w_3)$.

2.8.2 Inhomogene lineare Gleichungssysteme

Seien nun $A \in \mathbb{R}^{m \times n}$ und $b \in \mathbb{R}^m$ ein Spaltenvektor mit $b \neq 0$ (d.h. mindestens eine Komponente von b ist ungleich 0). Wir betrachten das lineare inhomogene Gleichungssystem

$$Ax = b.$$

Die Lösungsmenge $\text{Lös}(A, b) = \{x \in \mathbb{R}^n | Ax = b\}$ ist für $b \neq 0$ kein Untervektorraum des $\mathbb{R}^n$.

Beispiel 2.8.5 (Geraden im $\mathbb{R}^2$).

In $\mathbb{R}^2$ ist $\text{Lös}(A, b) = \{x \in \mathbb{R}^2 \mid a_1x_1 + a_2x_2 = b\}$ eine Geradengleichung. Die Gerade geht für $b \neq 0$ nicht durch den Ursprung, sondern entsteht durch *Parallelverschiebung* einer Ursprungsgeraden. Die Gleichung $Ax = 0$ heisst *zugehöriges homogenes Gleichungssystem*.

Definition 2.8.6 (Affiner Unterraum).

Eine Teilmenge X eines $\mathbb{R}$ -Vektorraumes V heisst **affiner Unterraum**, falls es ein $v \in V$ und einen Untervektorraum $L \subset V$ gibt, so dass

$$X = v + L$$

mit $v + L := \{w \in V \mid \exists l \in L \text{ mit } w = v + l\}$. Wir bezeichnen auch die leere Menge $\emptyset$ als affinen Unterraum. Affine Unterräume des $\mathbb{R}^n$ sind Punkte, Geraden, Ebenen etc.

Lemma 2.8.7. (Das Urbild eines Punktes bezüglich einer linearen Abbildung ist ein affiner Unterraum.)

Sei $F : V \rightarrow W$ eine lineare Abbildung. Dann ist für jedes $w \in W$ das Urbild $F^{-1}(w) \subset V$ ein affiner Unterraum. Ist $F^{-1}(w) \neq \emptyset$ und $v \in F^{-1}(w)$ beliebig, so ist

$$F^{-1}(w) = v + \text{Ker}(F). \quad (2.8)$$

Beweis: Im Fall $F^{-1}(w) = \emptyset$ ist nichts zu zeigen. Sei also $v \in F^{-1}(w)$. Für ein beliebiges $u \in F^{-1}(w)$ folgt wegen $F(u - v) = F(u) - F(v) = w - w = 0$, dass $u - v \in \text{Ker}(F)$ und somit $u \in v + \text{Ker}(F)$

Ist andererseits $u = v + v' \in v + \text{Ker}(F)$, dann gilt

$$F(u) = F(v) + F(v') = w + 0 = w,$$

also $u \in F^{-1}(w)$. Damit ist die Gleichheit der beiden Mengen in (2.8) gezeigt. $\square$

Aus Lemma 2.8.7 folgt sofort die analoge Aussage für lineare Gleichungssysteme, wenn wir $F : \mathbb{R}^n \rightarrow \mathbb{R}^m, x \mapsto Ax$ setzen:

Satz 2.8.8. (Die Lösungsmenge eines linearen Gleichungssystems ist ein affiner Unterraum.)

Sei $A \in \mathbb{R}^{m \times n}$ und $b \in \mathbb{R}^m$. Wir betrachten zu

$$Ax = b$$

die Lösungsmenge $\text{Lös}(A, b) = \{x \in \mathbb{R}^n \mid Ax = b\}$ und $\text{Ker}(A) = \{x \in \mathbb{R}^n \mid Ax = 0\}$. Ist $\text{Lös}(A, b) \neq \emptyset$ und $v \in \text{Lös}(A, b)$ beliebig (also $Av = b$), so ist $\text{Lös}(A, b) = v + \text{Ker}(A)$.

Merke: Die allgemeine Lösung $\text{Lös}(A, b)$ eines inhomogenen linearen Gleichungssystems erhält man durch Addition einer speziellen Lösung v mit $Av = b$ und der allgemeinen Lösung des homogenen Gleichungssystems, $\text{Ker}(A)$.

Die erweiterte Koeffizientenmatrix

Wir führen nun ein nützliches Hilfsmittel zur praktischen Berechnung der Lösung eines inhomogenen linearen Gleichungssystems ein: die **erweiterte Koeffizientenmatrix**. Dies ist die Matrix $(A, b) \in \mathbb{R}^{m \times (n+1)}$ mit

$$(A, b) := \left(\begin{array}{ccc|c} a_{11} & \dots & a_{1n} & b_1 \\ \vdots & & \vdots & \vdots \\ a_{m1} & \dots & a_{mn} & b_n \end{array} \right).$$

Satz 2.8.9 (Bedingung für Lösbarkeit).

Der Lösungsraum $\text{Lös}(A, b)$ des inhomogenen Gleichungssystems $Ax = b$ ist genau dann nicht leer, wenn $\text{Rang} A = \text{Rang}(A, b)$.

Definition 2.8.10 (Universelle und eindeutige Lösbarkeit).

Für festes $A \in \mathbb{R}^{m \times n}$ heisst das Gleichungssystem $Ax = b$ **universell lösbar**, falls es für jedes $b \in \mathbb{R}^n$ mindestens eine Lösung hat.

Ist b gegeben und hat die Lösungsmenge $\text{Lös}(A, b)$ genau ein Element, so heisst das Gleichungssystem **eindeutig lösbar**.

Merke:

1. (a) $Ax = b$ ist universell lösbar $\Leftrightarrow \text{Rang} A = m$.
2. (b) $Ax = b$ ist eindeutig lösbar $\Leftrightarrow \text{Rang}(A) = \text{Rang}(A, b) = n$.

2.8.3 Praktisches Lösungsverfahren

Starte mit der erweiterten Koeffizientenmatrix $A' = (A, b)$. Bringe (A, b) auf Zeilenstufenform (mit elementaren Zeilenumformungen)

$$\left(\begin{array}{cccc|c} 0 & b_{1j_1} & \dots & & c_1 \\ & 0 & b_{2j_2} & \dots & \vdots \\ & & & \ddots & \\ & 0 & & 0 & b_{rj_r} & \dots & c_r \\ & & & & & & c_{r+1} \\ & & & & & & \vdots \end{array} \right) = (B, c).$$

Es ist $b_{1j_1} \neq 0, \dots, b_{rj_r} \neq 0$. Dann ist $\text{Rang} A = r$. Wegen $\text{Rang}(A, b) = \text{Rang}(B, c)$ ist $\text{Rang}(A, b) = \text{Rang}(A) \Leftrightarrow c_{r+1} = \dots = c_m = 0$.

Denn: Nach eventueller Zeilenvertauschung wäre o.B.d.A. (ohne Beschränkung der Allgemeinheit) $c_{r+1} \neq 0$ und $0x_1 + \dots + 0x_n = c_{r+1}$ ist unlösbar! Sei also $c_{r+1} = \dots = c_m = 0$. Dann ist $\text{Lös}(A, b) \neq \emptyset$.

- (a) Wir müssen zuerst eine *spezielle* Lösung bestimmen.

- (a1) Die Unbestimmten x_j mit $j \notin \{j_1, \dots, j_r\}$ sind wieder freie Parameter. O.b.d.A. sei wieder $j_1 = 1, \dots, j_r = r$.
- (a2) Wir setzen $x_{r+1} = \dots = x_n = 0$
- (a3) Aus der r -ten Zeile von (B, c) erhält man

$$b_{rr}x_r = c_r,$$

also ist x_r bestimmt.

- (a4) Entsprechend erhält man $x_{r-1}, \dots, x_1$, also insgesamt eine spezielle Lösung

$$v = (x_1, \dots, x_r, 0, \dots, 0)^T$$

mit $Av = b$. Hier verwenden wir die Tatsache, dass eine Lösung von $Bx = c$, wobei (B, c) aus (A, b) durch elementare Zeilenumformung entsteht, auch Lösung von $Ax = b$ ist.

- (b) Nun ist nach Satz 2.8.8 nur noch die allgemeine Lösung des zugehörigen linearen homogenen Gleichungssystems

$$Ax = 0$$

zu bestimmen, denn $\text{Lös}(A, b) = v + \text{Ker}(A)$.

Beispiel 2.8.11. $A \in \mathbb{R}^{3 \times 4}$:

$$\begin{array}{cccccccl} x_1 & -2x_2 & +x_3 & & & = & 1 \\ x_1 & -2x_2 & & -x_4 & & = & 2 \\ & & x_3 & +x_4 & & = & -1 \end{array}$$

Wir bilden die erweiterte Koeffizientenmatrix:

$$\left(\begin{array}{cccc|c} 1 & -2 & 1 & 0 & 1 \\ 1 & -2 & 0 & -1 & 2 \\ 0 & 0 & 1 & 1 & -1 \end{array} \right) = (A, b),$$

bringen sie durch elementare Zeilenumformungen auf Zeilenstufenform

$$\left(\begin{array}{cccc|c} 1 & -2 & 1 & 0 & 1 \\ 0 & 0 & 1 & 1 & -1 \\ 0 & 0 & 0 & 0 & 0 \end{array} \right) = (B, c)$$

und erhalten das reduzierte Gleichungssystem:

$$\begin{array}{cccccccl} x_1 & -2x_2 & +x_3 & & & = & 1 \\ & & & x_3 & +x_4 & = & -1. \end{array}$$

Wegen $r = \text{Rang} A = \text{Rang}(A, b) = 2$ ist das Gleichungssystem lösbar.

$$\dim \text{Ker}(A) = n - r = 4 - 2 = 2, j_1 = 1, j_2 = 3.$$

Setze $x_2 = x_4 = 0$, und somit $x_3 = -1$ $x_1 + x_3 = 1 \Rightarrow x_1 = 1 - x_3 = 1 + 1 = 2$, also erhalten wir die spezielle Lösung $v = (2, 0, -1, 0)^T$. Die allgemeine Lösung von $Ax = 0$, mit $x_2 = \lambda_1$ und $x_4 = \lambda_2$ ist

$$\begin{aligned}x_3 &= -\lambda_2 \\x_1 &= 2\lambda_1 + \lambda_2,\end{aligned}$$

und somit gilt $x = (2\lambda_1 + \lambda_2, \lambda_1, -\lambda_2, \lambda_2)^T$. Mit $\lambda_1 = 1, \lambda_2 = 0$ erhalten wir $w_1 = (2, 1, 0, 0)^T$ und mit $\lambda_1 = 0, \lambda_2 = 1$ $w_2 = (1, 0, -1, 1)^T$. Wir erhalten also als allgemeine Lösung:

$$\text{Lös}(A, b) = \begin{pmatrix} 2 \\ 0 \\ -1 \\ 0 \end{pmatrix} + \text{Spann} \left(\begin{pmatrix} 2 \\ 1 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \\ -1 \\ 1 \end{pmatrix} \right).$$

Kapitel 3

Lineare Algebra II

In diesem Kapitel werden wir lernen, Vektorräume unabhängig von einer speziellen Basis zu betrachten. Dies erlaubt uns ein ganz neues, tieferes Verständnis von Matrizen und linearen Abbildungen zu gewinnen, mit dem man z.B. Phänomene wie Resonanz oder Abklingverhalten bei dynamischen Systemen erklären kann. Wir betrachten insbesondere für einen n -dimensionalen reellen Vektorraum V lineare Abbildungen von V nach V . Solche Abbildungen nennt man **Endomorphismen**. Da Endomorphismen Vektoren aus einem Vektorraum V wieder auf Vektoren aus V abbilden, können sie wiederholt angewendet werden.

In der Matrixdarstellung haben Endomorphismen die Form einer quadratischen Matrix, und wir werden uns in diesem Kapitel fast nur mit quadratischen Matrizen beschäftigen, außer in Kapitel 3.3.1.

3.1 Determinanten

Wir beginnen mit einer wichtigen Zahl, die man zu jeder quadratischen Matrix berechnen kann, der Determinante.

3.1.1 Determinante einer (2×2) -Matrix

Wir betrachten eine lineare Gleichung in $\mathbb{R}$: $a \cdot x = b$ mit $a \neq 0$. Die Lösung ist offensichtlich $x = \frac{b}{a}$. Wie wir sehen, ist sie als Ausdruck von a und b explizit darstellbar.

Fragen: Gilt diese letzte Beobachtung über die Darstellbarkeit der Lösung, falls eine solche existiert und eindeutig ist, auch für Gleichungssysteme (lineare Gleichungen in $\mathbb{R}^n$):

$$Ax = b \quad \text{mit } A \in \mathbb{R}^{n \times n}, b \in \mathbb{R}^n. \quad (3.1)$$

Und was sind Bedingungen für die Lösbarkeit von (3.1)?

Beispiel 3.1.1 (Determinante einer (2×2) -Matrix).

Sei $n = 2$. Ein lineares Gleichungssystem in $\mathbb{R}^2$ mit zwei Gleichungen hat die allgemeine Form

$$\begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} b_1 \\ b_2 \end{pmatrix}. \quad (3.2)$$

Falls $a_{11}a_{22} - a_{21}a_{12} \neq 0$, dann ist die eindeutige Lösung

$$x_1 = \frac{b_1a_{22} - b_2a_{12}}{a_{11}a_{22} - a_{21}a_{12}}, \quad x_2 = \frac{a_{11}b_2 - a_{21}b_1}{a_{11}a_{22} - a_{21}a_{12}},$$

wie man z.B. mit Hilfe des Gauß-Algorithmus herleiten kann.

Wir definieren

$$\begin{aligned} \det \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} &:= \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} \\ &:= a_{11}a_{22} - a_{21}a_{12}. \end{aligned} \quad (3.3)$$

Mit dieser Notation können wir die Lösung (3.2) wie folgt schreiben:

$$x_1 = \frac{\begin{vmatrix} b_1 & a_{12} \\ b_2 & a_{22} \end{vmatrix}}{\begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix}}, \quad x_2 = \frac{\begin{vmatrix} a_{11} & b_1 \\ a_{21} & b_2 \end{vmatrix}}{\begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix}} \quad (3.4)$$

Wir bezeichnen $\det(A)$ als die **Determinante** von A .

Bemerkung 3.1.2 (Determinanten von $(n \times n)$ -Matrizen).

Die Determinante ist auch für größere quadratische Matrizen definiert, wie wir bald sehen werden, und es gibt ein ähnliches Lösungsverfahren wie das vorgestellte auch für $n \geq 3$, die sogenannte Cramersche Regel. Dieses Verfahren hat für praktische Berechnungen aber keine Relevanz. Determinanten von allgemeinen $(n \times n)$ -Matrizen werden trotzdem für die weitere Vorlesung wichtig sein, z.B. zur Definition des charakteristischen Polynoms einer Matrix (s. Definition 3.2.7) und bei der Integration im $\mathbb{R}^2$ mit Polarkoordinaten.

Wir beobachten folgende Eigenschaften der Determinante von (2×2) -Matrizen (3.3):

1. Notwendige und hinreichende Bedingung dafür, daß $Ax = b$ für jedes b eindeutig lösbar ist, d.h. für die Invertierbarkeit von A , ist $\det A \neq 0$.
2. Der Ausdruck $\det A = a_{11}a_{22} - a_{21}a_{12}$ ist der orientierte (mit Vorzeichen) Flächeninhalt des von den Zeilenvektoren $v_1 = (a_{11}, a_{12})$ und $v_2 = (a_{21}, a_{22})$ aufgespannten Parallelogramms (siehe Abbildung 3.1) Dank dieser geometrischen Deutung erkennen wir sofort folgende leicht nachzurechnende Eigenschaften der Determinante (3.3):
 - (a) Das System (v_1, v_2) ist genau dann linear abhängig, wenn das entsprechende Parallelogramm entartet ist, d.h. die Fläche Null hat.
 - (b) Bei Vertauschung der beiden Zeilen ändert sich das Vorzeichen der Determinante, da das entsprechende Parallelogramm seine Orientierung wechselt:

$$\det \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = -\det \begin{pmatrix} v_2 \\ v_1 \end{pmatrix}. \quad (3.5)$$

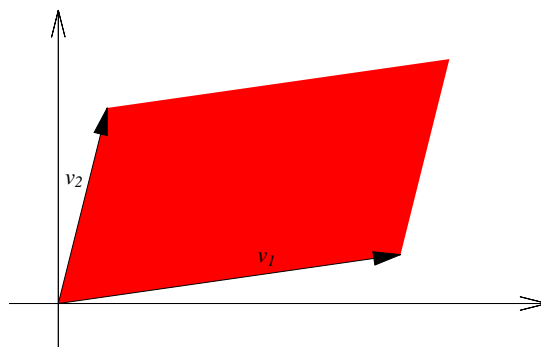


Abbildung 3.1: Die Determinante entspricht dem orientierten Flächeninhalt des von v_1 und v_2 aufgespannten Parallelogramms.

- (c) Die Determinante ändert sich nicht, wenn man ein skalares Vielfaches einer Zeile zu einer anderen Zeile addiert, da das Volumen sich bei Scherung nicht ändert (vgl. Abbildung 3.2):

$$\det \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = \det \begin{pmatrix} v_1 \\ v_2 + \lambda \cdot v_1 \end{pmatrix}. \quad (3.6)$$

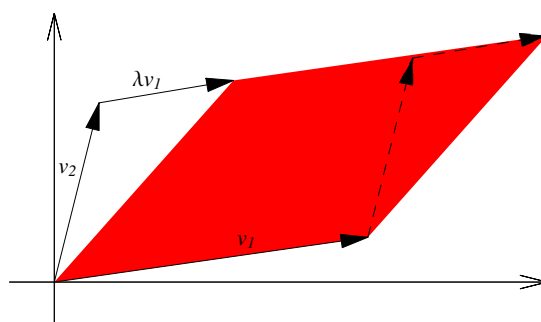


Abbildung 3.2: Die Fläche des Parallelogramms bleibt gleich, wenn v_2 durch $v_2 + \lambda v_1$ ersetzt wird.

- (d) Multipliziert man eine Zeile mit $\lambda \in \mathbb{R}$, so multipliziert sich auch die Determinante mit λ . Für $\lambda > 0$ entspricht dies der Streckung des Parallelogramms um einen Faktor λ in Richtung des entsprechenden Zeilenvektors:

$$\det \begin{pmatrix} \lambda \cdot v_1 \\ v_2 \end{pmatrix} = \lambda \cdot \det \begin{pmatrix} v_1 \\ v_2 \end{pmatrix}. \quad (3.7)$$

- (e) Unterscheiden sich zwei (2×2) -Matrizen A und B in nur einer Zeile (mit Zeilenindex i), so ist die Summe ihrer Determinanten gleich der Determinante der Matrix C , deren i -te Zeile gleich der Summe der i -ten Zeilen von A und B ist und die in der anderen Zeile mit A und B übereinstimmt. Wie man nämlich in Abbildung 3.3

für das Beispiel $i = 1$ erkennt, hat das Parallelogramm der Matrix $C = \begin{pmatrix} v_1 + \tilde{v}_1 \\ v_2 \end{pmatrix}$ den gleichen Flächeninhalt wie die beiden Matrizen $A = \begin{pmatrix} v_1 \\ v_2 \end{pmatrix}$ und $B = \begin{pmatrix} \tilde{v}_1 \\ v_2 \end{pmatrix}$ entsprechenden Parallelogramme. Dazu legen wir diese an jeweiligen Kanten aneinander, die den identischen Zeilenvektoren entsprechen:

$$\det \begin{pmatrix} v_1 + \tilde{v}_1 \\ v_2 \end{pmatrix} = \det \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} + \det \begin{pmatrix} \tilde{v}_1 \\ v_2 \end{pmatrix}. \quad (3.8)$$

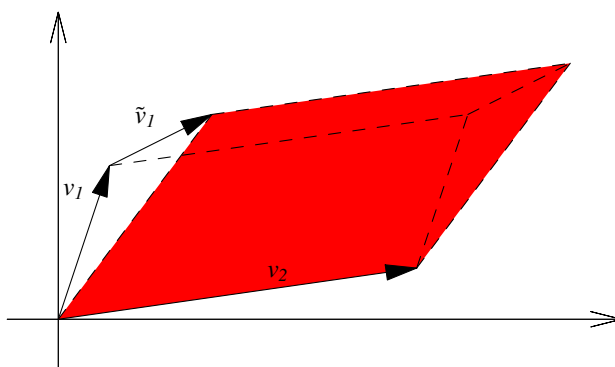


Abbildung 3.3: Die Summe von zwei Parallelogrammen mit gemeinsamer Kante

Die Gleichungen (3.7) und (3.8) bedeuten, dass die Determinante linear in jeder Zeile ist.

3.1.2 *Permutationen

Für eine explizite Darstellung der Determinante einer $(n \times n)$ -Matrix benötigen wir einige Begriffe aus der Gruppentheorie.

Definition 3.1.3 (symmetrische Gruppe S_n).

Für jede natürliche Zahl $n > 0$ sei S_n die **symmetrische Gruppe** von $\{1, \dots, n\}$, d.h. die Menge aller bijektiven Abbildungen

$$\sigma : \{1, \dots, n\} \rightarrow \{1, \dots, n\}.$$

Die Elemente von S_n heißen **Permutationen**. Eine Permutation $\sigma \in S_n$ lässt sich folgendermaßen darstellen:

$$\sigma = \begin{bmatrix} 1 & 2 & 3 & \dots & n \\ \sigma(1) & \sigma(2) & \sigma(3) & \dots & \sigma(n) \end{bmatrix}.$$

Beispiel 3.1.4 (für eine Permutation).

Ein Beispiel wäre z.B. die folgende Permutation $\sigma \in S_4$:

$$\begin{array}{ccc}
 1 & \nearrow & 1 & = & \sigma(2) \\
 2 & \searrow & 2 & = & \sigma(3) \\
 3 & \nearrow & 3 & = & \sigma(1) \\
 \\
 4 & \longrightarrow & 4 & = & \sigma(4)
 \end{array}$$

mit der Permutationstafel: $\begin{bmatrix} 1 & 2 & 3 & 4 \\ 3 & 1 & 2 & 4 \end{bmatrix}$.

Für $\tau, \sigma \in S_n$ gilt

$$\tau \circ \sigma = \begin{bmatrix} 1 & \dots & n \\ \tau(1) & \dots & \tau(n) \end{bmatrix} \circ \begin{bmatrix} 1 & \dots & n \\ \sigma(1) & \dots & \sigma(n) \end{bmatrix} \quad (3.9)$$

$$= \begin{bmatrix} 1 & \dots & n \\ \tau(\sigma(1)) & \dots & \tau(\sigma(n)) \end{bmatrix}. \quad (3.10)$$

Mit „ $\circ$ “ ist die Gruppen-Verknüpfung gemeint.

Beispiel 3.1.5 (Nicht kommutierende Permutationen).

Es gilt

$$\begin{array}{l}
 \begin{bmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \end{bmatrix} \circ \begin{bmatrix} 1 & 2 & 3 \\ 1 & 3 & 2 \end{bmatrix} = \begin{bmatrix} 1 & 2 & 3 \\ 2 & 1 & 3 \end{bmatrix}, \\
 \text{aber} \quad \begin{bmatrix} 1 & 2 & 3 \\ 1 & 3 & 2 \end{bmatrix} \circ \begin{bmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \end{bmatrix} = \begin{bmatrix} 1 & 2 & 3 \\ 3 & 2 & 1 \end{bmatrix}.
 \end{array}$$

Die Gruppe S_n ist für $n \geq 3$ nicht kommutativ!

Bemerkung 3.1.6. Die Gruppe S_n hat genau $n!$ Elemente.

Wir führen noch Funktion auf der Menge der Permutationen ein, die wir für eine explizite Formel der Determinante benötigen.

***Definition 3.1.7** (Signum-Funktion für Permutationen, Fehlstand).

Das **Signum** einer Permutation σ ist definiert durch

$$\text{sign}(\sigma) := \begin{cases} +1 & : \sigma \text{ hat gerade Anzahl Fehlstände,} \\ -1 & : \sigma \text{ hat ungerade Anzahl Fehlstände.} \end{cases}$$

Ein **Fehlstand** von $\sigma \in S_n$ ist ein Paar $i, j \in \{1, \dots, n\}$ mit $i < j$ und $\sigma(i) > \sigma(j)$.

3.1.3 Eigenschaften der Determinante

In (3.3) haben wir schon für jede (2×2) -Matrix deren Determinante durch eine explizite Formel definiert und in Abschnitt 3.1.1 deren Eigenschaften beobachtet. Nun gehen wir umgekehrt vor. Wir definieren jetzt Determinanten allgemein für $(n \times n)$ -Matrizen durch ihre Eigenschaften und zeigen anschließend die Existenz und Eindeutigkeit der Determinante und geben auch eine explizite Formel für sie an.

Ist A eine n -zeilige quadratische Matrix, so werden im folgenden mit $a_1, \dots, a_n$ die Zeilenvektoren von A bezeichnet. Es ist also

$$A = \begin{pmatrix} a_1 \\ \vdots \\ a_n \end{pmatrix}. \quad (3.11)$$

Definition 3.1.8 (Determinante).

Eine **Determinante** ist eine Abbildung

$$\det : \mathbb{R}^{n \times n} \rightarrow \mathbb{R},$$

für alle $n > 0$, mit folgenden Eigenschaften:

1. $\det$ ist linear in jeder Zeile.

Genauer: Ist $A \in \mathbb{R}^{n \times n}$ wie in (3.11) und $i \in \{1, \dots, n\}$, so gilt:

- (a) Ist $a_i = a'_i + a''_i$, so ist

$$\det \begin{pmatrix} \vdots \\ a_i \\ \vdots \end{pmatrix} = \det \begin{pmatrix} \vdots \\ a'_i \\ \vdots \end{pmatrix} + \det \begin{pmatrix} \vdots \\ a''_i \\ \vdots \end{pmatrix}$$

- (b) Ist $a_i = \lambda a'_i$, so ist

$$\det \begin{pmatrix} \vdots \\ a_i \\ \vdots \end{pmatrix} = \lambda \det \begin{pmatrix} \vdots \\ a'_i \\ \vdots \end{pmatrix}$$

Dabei stehen die Punkte jeweils für die Zeilenvektoren $a_1, \dots, a_{i-1}, a_{i+1}, \dots, a_n$.

2. $\det$ ist alternierend, d.h. hat A zwei gleiche Zeilen, so ist $\det A = 0$.
3. $\det$ ist normiert, d.h. $\det \mathbb{I}_n = 1$.

Satz 3.1.9 (Eigenschaften der Determinante).

Die Determinante $\det : \mathbb{R}^{n \times n} \rightarrow \mathbb{R}$ hat die folgenden weiteren Eigenschaften

1. Für alle $\lambda \in \mathbb{R}$ ist

$$\det(\lambda A) = \lambda^n \det A.$$

2. Gibt es ein i mit $a_i = (0, \dots, 0)$ so ist $\det A = 0$.

3. Entsteht B aus A durch eine Zeilenvertauschung, so ist $\det B = -\det A$, also

$$\det \begin{pmatrix} \vdots \\ a_j \\ \vdots \\ a_i \\ \vdots \end{pmatrix} = -\det \begin{pmatrix} \vdots \\ a_j \\ \vdots \\ a_i \\ \vdots \end{pmatrix}. \quad (3.12)$$

4. Ist $\lambda \in \mathbb{R}$ und entsteht B aus A durch Addition der λ -fachen j -ten Zeile zur i -ten Zeile ($i \neq j$), so ist $\det B = \det A$, also

$$\det \begin{pmatrix} \vdots \\ a_i + \lambda a_j \\ \vdots \\ a_j \\ \vdots \end{pmatrix} = \det \begin{pmatrix} \vdots \\ a_i \\ \vdots \\ a_j \\ \vdots \end{pmatrix}.$$

5. Ist A eine obere Dreiecksmatrix, i.e.

$$A = \begin{pmatrix} \lambda_1 & \dots & \\ & \ddots & \vdots \\ 0 & & \lambda_n \end{pmatrix},$$

wobei die Koeffizienten nur auf und oberhalb der Diagonalen von 0 verschiedene Werte annehmen können, so ist

$$\det A = \lambda_1 \cdot \lambda_2 \cdot \dots \cdot \lambda_n. \quad (3.13)$$

6. $\det A = 0$ ist gleichbedeutend damit, daß die Zeilenvektoren $a_1, \dots, a_n$ linear abhängig sind.
7. Ist $\det A \neq 0$ so ist A invertierbar.
8. Für $A, B \in \mathbb{R}^{n \times n}$ gilt der **Determinantenmultiplikationssatz**:

$$\det(A \cdot B) = \det(A) \cdot \det(B).$$

Insbesondere gilt für invertierbare Matrizen A :

$$\det(A^{-1}) = (\det A)^{-1}.$$

9. Es gilt

$$\det(A) = \det(A^T).$$

Daraus folgt, dass zu den Aussagen (3.), (4.) und (6.) über die *Zeilen* einer Matrix analoge Aussagen über die *Spalten* einer Matrix gelten.

Fundamental ist der folgender Satz.

***Satz 3.1.10** (Eindeutigkeit der Determinante).

Es gibt *genau eine* Determinante

$$\det : \mathbb{R}^{n \times n} \rightarrow \mathbb{R}, n > 0,$$

und zwar ist für $A = (a_{ij})_{\substack{1 \leq i \leq n \\ 1 \leq j \leq n}}$:

$$\det A = \sum_{\sigma \in S_n} \text{sign}(\sigma) \cdot a_{1\sigma(1)} \cdots a_{n\sigma(n)}.$$

Dabei haben wir die Signum Funktion verwendet (s. Definition 3.1.7).

Notation: Wir schreiben auch

$$\det \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \cdots & a_{nn} \end{pmatrix} =: \begin{vmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \cdots & a_{nn} \end{vmatrix}.$$

Beispiel 3.1.11 (Determinanten von $(n \times n)$ -Matrizen für $n \in \{1, 2, 3\}$).

$$n = 1 : \quad \det(a) = a. \quad (3.14)$$

$$n = 2 : \quad \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11}a_{22} - \underbrace{a_{12}a_{21}}_{\text{Fehlstand (1, 2)}}. \quad (3.15)$$

$$n = 3 : \quad \begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} = a_{11}a_{22}a_{33} - \underbrace{a_{11}a_{23}a_{32}}_{(1)} - \underbrace{a_{12}a_{21}a_{33}}_{(2)} + \underbrace{a_{12}a_{23}a_{31}}_{(3)} \\ + \underbrace{a_{13}a_{21}a_{32}}_{(4)} - \underbrace{a_{13}a_{22}a_{31}}_{(5)} \quad (3.16)$$

In (3.16) treten folgende Fehlstände auf:

- (1) Fehlstand (2, 3).
- (2) Fehlstande(1, 2).
- (3) Fehlstände (1, 3) und (2, 3).
- (4) Fehlstände (2, 3) und (1, 2).
- (5) Fehlstände (1, 2), (1, 3) und (2, 3).

Wir bemerken noch, dass die Summe in (3.16) genau $3! = 6$ Summanden hat.

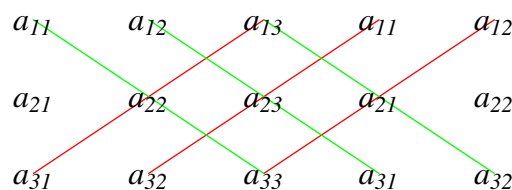


Abbildung 3.4: Illustration des Schemas von Sarrus

Man kann sich Formel (3.16) auch mit Hilfe des folgenden Schemas merken (nach Sarrus): Die Produkte längs der Hauptdiagonalen (nach rechts unten) haben positives Vorzeichen, solche längs der Nebendiagonalelemente haben negatives Vorzeichen.

3.1.4 Praktische Berechnung von Determinanten

Sei $A \in \mathbb{R}^{n \times n}$ gegeben. Durch Zeilenumformungen vom Typ U_2 und U_3 (vgl. 2.7.13) kann A auf Zeilenstufenform B gebracht werden. Mit Hilfe der Eigenschaften 3.1.8.1 und 3.1.8.2 der Determinanten in Definition 3.1.8 folgt dann

$$\det A = (-1)^k \det B,$$

wobei k die Anzahl der elementaren Umformung vom Typ U_3 ist. Nach Eigenschaft 5 in Satz 3.1.9 ist $\det B$ das Produkt der Diagonalelemente.

Beispiel 3.1.12 (Berechnung der Determinante einer (3×3) -Matrix).

Wir berechnen folgende Determinante mit Hilfe von elementaren Zeilenumformungen.

$$\begin{aligned} \begin{vmatrix} 0 & 1 & 3 \\ 3 & 2 & 1 \\ 1 & 1 & 0 \end{vmatrix} &= - \begin{vmatrix} 1 & 1 & 0 \\ 3 & 2 & 1 \\ 0 & 1 & 3 \end{vmatrix} = - \begin{vmatrix} 1 & 1 & 0 \\ 0 & -1 & 1 \\ 0 & 1 & 3 \end{vmatrix} \\ &= - \begin{vmatrix} 1 & 1 & 0 \\ 0 & -1 & 1 \\ 0 & 0 & 4 \end{vmatrix} = 4. \end{aligned}$$

Zur Kontrolle berechnen wir die Determinante auch noch mit der Regel von Sarrus:

$$\begin{aligned} \begin{vmatrix} 0 & 1 & 3 \\ 3 & 2 & 1 \\ 1 & 1 & 0 \end{vmatrix} &= 0 \cdot 2 \cdot 0 + 1 \cdot 1 \cdot 1 + 3 \cdot 3 \cdot 1 - 1 \cdot 2 \cdot 3 - 1 \cdot 1 \cdot 0 - 0 \cdot 3 \cdot 1 \\ &= 4. \end{aligned}$$

Beispiel 3.1.13 (Laplacescher Entwicklungssatz).

Ein anderes Verfahren, mit dem man Determinanten berechnen kann, spaltet die gegebene Matrix in kleinere Untermatrizen auf. Die Determinante wird hier nach einer Zeile (oder Spalte) entwickelt, d.h. man geht nacheinander die Elemente dieser Zeile (Spalte) durch, multipliziert sie jeweils mit der Determinante einer Untermatrix, und addiert sie dann mit wechselndem Vorzeichen auf. Um zu jedem Element die entsprechende Untermatrix zu erhalten, streicht man die Zeile und die Spalte, die dem jeweiligen Element entsprechen, und erhält aus den übriggebliebenen Matrixelementen wieder eine quadratische Matrix mit einer Dimension weniger, deren Determinante leichter zu berechnen ist.

Zur Illustration rechnen wir die Determinante aus dem obigen Beispiel noch einmal mit diesem Verfahren aus, wobei wir nach der ersten Zeile entwickeln:

$$\begin{vmatrix} 0 & 1 & 3 \\ 3 & 2 & 1 \\ 1 & 1 & 0 \end{vmatrix} = +0 \begin{vmatrix} 2 & 1 \\ 1 & 0 \end{vmatrix} - 1 \begin{vmatrix} 3 & 1 \\ 1 & 0 \end{vmatrix} + 3 \begin{vmatrix} 3 & 2 \\ 1 & 1 \end{vmatrix} = +0 \cdot (1) - 1 \cdot (-1) + 3 \cdot (3 - 2) = 4$$

Für die Vorzeichen bei der Summation der Beiträge jedes Elements der Zeile (bzw. Spalte), nach der wir entwickeln, gilt folgendes „Schachbrettmuster“:

$$\begin{array}{cccccc} + & - & + & - & + & \dots \\ - & + & - & + & - & \dots \\ + & - & + & - & + & \dots \\ - & + & - & + & - & \dots \\ + & - & + & - & + & \dots \\ \vdots & & & & & \end{array}$$

Als Übung könnte man die Determinante nach der zweiten Spalte berechnen. Welches Ergebnis erwarten Sie?

3.2 Eigenwerte und Eigenvektoren

Wir kommen nun auf ein wichtiges Konzept der linearen Algebra zu sprechen, nämlich zu Eigenwerten und des Eigenvektoren von Endomorphismen bzw. von quadratischen Matrizen. Zur Motivation betrachten wir ein Beispiel aus der Populationsdynamik.

Modell 1 (Lineares Populationsmodell mit einer Zustandsvariablen).

Sei $v^{(k)}$ die Anzahl der Paare (Männchen und Weibchen) von Kaninchen im Monat k ($k = 0, 1, 2, \dots$). Im Monat $k + 1$ hat jedes Paar Nachwuchs bekommen, und zwar genau c Paare (jeweils ein Männchen und ein Weibchen), wobei $c \in \{0, 1, 2, \dots\}$. Im Monat 0 gebe es genau a Paare ($a \in \{0, 1, \dots\}$). Wir erhalten also eine **Differenzengleichung mit Anfangsbedingung**:

$$\begin{cases} v^{(0)} & = & a & \text{Anfangsbedingung,} \\ v^{(k+1)} - v^{(k)} & = & c \cdot v^{(k)} & \text{Differenzengleichung} \end{cases}$$

Bemerkung 3.2.1. Modell 1 ist sehr simpel, da von einer konstanten Vermehrungsrate ausgegangen wird, ohne Rücksicht auf äußere Bedingungen wie z.B.: Gesamtzahl der Paare und Ressourcen, individuelle Eigenschaften der Kaninchen (Alter). Der Tod von Kaninchen wird auch nicht berücksichtigt. Wir betrachten aber zur Illustration absichtlich ein solch einfaches Modell.

Der *Zustand* des Systems zu einem bestimmten Zeitpunkt wird durch eine Zahl $\in \mathbb{R}$ (1-dim reeller Vektorraum) beschrieben. Der Übergang von einem Zustand (im Monat k) zum nächsten (im Monat $k + 1$) wird durch eine *lineare Abbildung* beschrieben:

$$v^{(k+1)} = (c + 1)v^{(k)}. \quad (3.17)$$

Wir finden leicht eine explizite Darstellung für $v^{(k)}$ (der Lösung des Anfangswertproblems) für allgemeines $k \in \mathbb{N}$:

$$v^{(k)} = (c + 1)^k \cdot a \quad (3.18)$$

Dabei können wir $(c + 1)^k$ als die k -malige Anwendung der linearen Multiplikation mit der Zahl $c + 1$ verstehen. Für $a > 0$ und $c > 0$ erhalten wir **exponentielles Wachstum**.

Modell 2 (Altersstrukturierte Kaninchenpopulationen).

Wir ändern Modell 1 leicht ab. Neugeborene Kaninchen können sich nicht in ihrem ersten Lebensmonat fortpflanzen, sondern erst ab dem zweiten. Wir beschreiben den Zustand des Systems im k -ten Monat durch den Vektor

$$v^{(k)} = \begin{pmatrix} v_1^{(k)} \\ v_2^{(k)} \end{pmatrix} \in \mathbb{N}^2 \subset \mathbb{R}^2,$$

wobei $v_1^{(k)}$ die Zahl der im Monat k *neugeborenen (jungen)* Paare ist und $v_2^{(k)}$ die Zahl der *alten* Paare (älter als ein Monat). Z.B. entspricht ein junges Paar dem Vektor $\begin{pmatrix} 1 \\ 0 \end{pmatrix}$.

Die Anfangsbedingung sei $\begin{pmatrix} v_1^{(0)} \\ v_2^{(0)} \end{pmatrix} = a = \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \in \mathbb{N}^2 \subset \mathbb{R}^2$. Jedes alte Paar zeugt jeden Monat c Paare. Wir haben also einen Übergang

$$\begin{pmatrix} 0 \\ 1 \end{pmatrix} \mapsto \begin{pmatrix} 0 \\ 1 \end{pmatrix} + c \cdot \begin{pmatrix} 1 \\ 0 \end{pmatrix}$$

von einem Monat auf den nächsten.

Junge Paare zeugen noch keinen Nachwuchs, werden aber in einem Zeitschritt (1 Monat) alt, also

$$\begin{pmatrix} 1 \\ 0 \end{pmatrix} \mapsto \begin{pmatrix} 0 \\ 1 \end{pmatrix}.$$

Wir erhalten die Rekursionsformel

$$\begin{aligned} v^{(k+1)} = \begin{pmatrix} v_1^{k+1} \\ v_2^{k+1} \end{pmatrix} &= \begin{pmatrix} 0 & c \\ 1 & 1 \end{pmatrix} \begin{pmatrix} v_1^{(k)} \\ v_2^{(k)} \end{pmatrix} \\ &= A \cdot v^{(k)}. \end{aligned} \quad (3.19)$$

Beispiel:

$$\begin{aligned} c &= 1, \quad a = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \\ v^{(0)} &= \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad v^{(1)} = \begin{pmatrix} 0 \\ 1 \end{pmatrix}, \quad v^{(2)} = \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \\ v^{(3)} &= \begin{pmatrix} 1 \\ 2 \end{pmatrix}, \quad v^{(4)} = \begin{pmatrix} 2 \\ 3 \end{pmatrix}, \dots \end{aligned}$$

Wir interessieren uns für eine explizite Darstellung von $v^{(k)}$, analog zu (3.18). Anhand dieser könnten wir z.B. untersuchen, ob das Wachstum der Gesamtpopulation auch exponentiell ist, und wenn ja, wie groß die Wachstumsrate ist.

Offensichtlich erhalten wir (durch Abspulen der Rekursionsgleichung (3.19))

$$\begin{aligned} \begin{pmatrix} v_1^{(k)} \\ v_2^{(k)} \end{pmatrix} &= \begin{pmatrix} 0 & c \\ 1 & 1 \end{pmatrix}^k \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \\ \Leftrightarrow \quad v^{(k)} &= A^k \cdot a. \end{aligned}$$

Wir wollen also für beliebiges k den Vektor $A^k \cdot a$ berechnen.

Allgemeine Frage: Wie „berechnet“ man für $a \in \mathbb{R}^m$, $A \in \mathbb{R}^{m \times m}$ und $k \in \mathbb{N}$ den Vektor $A^k a$?

Antwort: Das hängt davon ab, was mit „berechnen“ gemeint ist:

1. Für die ersten k Monate (wenn k ist nicht allzu groß ist), kann man $v^{(k)}$ per Hand oder mit dem Computer ausrechnen und grafisch darstellen, wie z.B. in Abbildung 3.5.
2. Wir sind aber auch an qualitativen Aussagen, z.B. dem Verhalten der Folge (Konvergenz, Divergenz) interessiert. Dazu wäre eine explizite Darstellung von $v^{(k)}$ analog zu (3.18) nützlich.

Unsere Aufgabe ist also: Berechne $A^k a = \underbrace{A \dots \cdot (A(Aa))}_{k\text{-mal}}$. Dazu müssen wir etwas weiter ausholen.

Als Heuristik verwenden wir das „Was wäre schön?“-Prinzip, d.h. wir überlegen uns, für welche a die Berechnung besonders einfach ist: Wenn für a gilt, dass $A \cdot a = \lambda \cdot a$, mit einem $\lambda \in \mathbb{R}$ oder $\lambda \in \mathbb{C}$, dann folgt daraus:

$$\begin{aligned} A^0 a &= a \\ A^1 a &= \lambda a \\ A^2 a &= A(Aa) = A(\lambda a) = \lambda \cdot Aa = \lambda^2 a \\ &\vdots \\ A^k a &= \lambda^k a. \end{aligned}$$

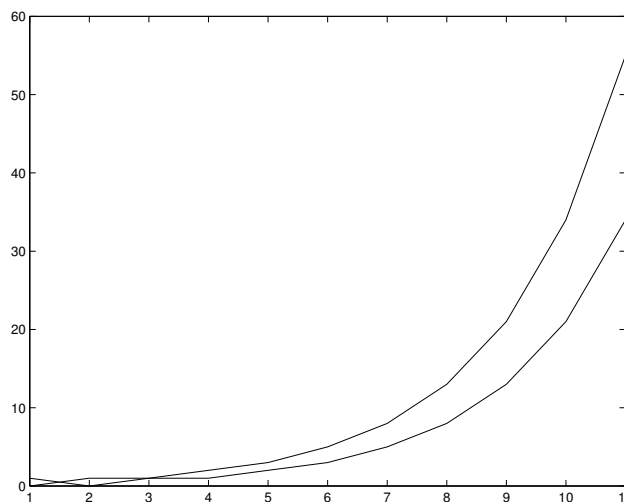


Abbildung 3.5: Die Kaninchenpopulation für die ersten 11 Monate, startend mit einem jungen Paar ($a = (1, 0)$), für die Vermehrungsrate $c = 1$.

Es gibt in der Tat solche Vektoren. Man nennt sie **Eigenvektoren** von A , und die entsprechende Zahl λ nennt man **Eigenwert**. Für Eigenvektoren von A ist die Multiplikation mit A^k also sehr einfach. Die Iteration erfolgt dann so leicht wie in Modell 1, einfach durch Potenzieren des Eigenwerts. Aber wie findet man Eigenvektoren und Eigenwerte?

Eine notwendige und hinreichende Bedingung dafür, dass $\lambda \in \mathbb{C}$ ein Eigenwert von $A \in \mathbb{R}^{n \times n}$ ist, ist die Existenz eines Eigenvektors $v \in \mathbb{R}^n \setminus \{0\}$ mit

$$\begin{aligned} Av &= \lambda v \\ \Leftrightarrow (A - \lambda \mathbb{I}_n)v &= 0, \end{aligned}$$

d.h. die Matrix $A - \lambda \mathbb{I}_n$, aufgefasst als lineare Abbildung des $\mathbb{C}^n$, muß einen nicht-trivialen Kern haben:

$$\text{Kern}(A - \lambda \mathbb{I}_n) \neq \{0\}.$$

Notwendige und hinreichende Bedingung hierfür ist

$$\det(A - \lambda \mathbb{I}_n) = 0.$$

Berechnung der Eigenwerte: Für unser Beispiel berechnen wir:

$$\begin{aligned} \det \left(\begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix} - \begin{pmatrix} \lambda & 0 \\ 0 & \lambda \end{pmatrix} \right) &= \det \begin{pmatrix} -\lambda & 1 \\ 1 & 1 - \lambda \end{pmatrix} \\ &= \lambda(\lambda - 1) - 1 \\ &= \lambda^2 - \lambda - 1 \\ &\stackrel{!}{=} 0. \end{aligned}$$

Die Lösungen dieser quadratischen Gleichung sind:

$$\begin{aligned}\lambda_1 &= \frac{1 - \sqrt{5}}{2} \approx -0.68034\dots \\ \lambda_2 &= \frac{1 + \sqrt{5}}{2} \approx 1.618\dots\end{aligned}$$

Bemerkung 3.2.2 (Goldener Schnitt).

Die Zahl $\tau := \frac{2}{1+\sqrt{5}} \approx 0.618\dots$ heißt **goldener Schnitt** und hat viele Menschen über die Jahrhunderte stark fasziniert. Der goldene Schnitt spielt u.a. in den bildenden Künsten und der Phyllotaxis eine große Rolle. Er erfüllt die einfache Gleichung

$$\tau = \frac{1}{1 + \tau}$$

und bezeichnet damit z.B. das Verhältnis zweier Längen a und b , die sich zueinander so verhalten, wie die längere der beiden zur gemeinsamen Summe: Falls $b > a$ dann folgt aus $\frac{a}{b} = \frac{b}{a+b}$ also, dass $\frac{a}{b} = \tau$. Bei den Kaninchenpopulationen kommt dieser Zusammenhang daher, dass das Verhältnis zwischen jungen und alten Kaninchen gegen τ konvergiert. Die Zahl der jungen zu der der alten Kaninchen verhält sich so wie die Zahl der jungen Kaninchen der nächsten Generation (die der Zahl der „alten“ alten entspricht, die Junge bekommen konnten) zu der der alten Kaninchen der nächsten Generation (die der Zahl der alten und jungen zusammen entspricht).

Berechnung der Eigenvektoren: Zu jedem λ_i berechnen wir einen Eigenvektor $w^{(i)} = \begin{pmatrix} w_1^{(i)} \\ w_2^{(i)} \end{pmatrix}$.

Zu $\lambda_1 = \frac{1-\sqrt{5}}{2}$: Bestimme den Kern $(A - \lambda_1 I_2)$, d.h. löse in $\mathbb{C}^2$ das lineare Gleichungssystem:

$$\left(\begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix} - \lambda_1 \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \right) \begin{pmatrix} w_1^{(1)} \\ w_2^{(1)} \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}. \quad (3.20)$$

Die Rechnung per Hand oder mit dem Computer ergibt den Eigenraum zu λ_1 , d.h. die Menge aller Lösungen zu (3.20).

$$\begin{aligned}E_{\lambda_1} &:= \text{Ker}(A - \lambda_1 I_2) \\ &= \text{Spann} \left\{ \begin{pmatrix} \frac{-1-\sqrt{5}}{2} \\ 1 \end{pmatrix} \right\}.\end{aligned}$$

Wir wählen

$$w^{(1)} = \begin{pmatrix} \left(\frac{-1-\sqrt{5}}{2} \right) \\ 1 \end{pmatrix}. \quad (3.21)$$

Wir berechnen ebenso zu $\lambda_2 = \frac{1+\sqrt{5}}{2}$:

$$E_\lambda = \text{Spann} \left\{ \begin{pmatrix} \frac{-1+\sqrt{5}}{2} \\ 1 \end{pmatrix} \right\}$$

und wählen

$$w^{(2)} = \begin{pmatrix} \frac{-1+\sqrt{5}}{2} \\ 1 \end{pmatrix}. \quad (3.22)$$

Berechnung von $A^k a$ für beliebige Vektoren $a \in \mathbb{R}^2$: Es gilt

$$A^k w^{(i)} = \lambda_i^k \cdot w^{(i)} \quad \text{für } i \in \{1, 2\}$$

und somit für jede Linearkombination $y_1 w^{(1)} + y_2 w^{(2)}$:

$$\text{und } A^k (y_1 w^{(1)} + y_2 w^{(2)}) = \lambda_1^k w^{(1)} + \lambda_2^k w^{(2)}.$$

Beobachtung: Das System $(w^{(1)}, w^{(2)})$ ist eine Basis des $\mathbb{R}^2$, denn eine Linearkombination $y_1 w^{(1)} + y_2 w^{(2)}$ ist genau dann gleich 0, wenn $y_1 = y_2 = 0$, da die Matrix $\begin{pmatrix} w_1^{(1)} & w_1^{(2)} \\ w_2^{(1)} & w_2^{(2)} \end{pmatrix}$ regulär ist (vgl. Definition 2.7.10) wegen

$$\det \begin{pmatrix} w_1^{(1)} & w_1^{(2)} \\ w_2^{(1)} & w_2^{(2)} \end{pmatrix} = \frac{-1 - \sqrt{5}}{2} \cdot 1 - 1 \cdot \frac{-1 + \sqrt{5}}{2} = -\sqrt{5}.$$

Wir können also jeden Vektor $a \in \mathbb{R}^2$ eindeutig als Linearkombination von $w^{(1)}$ und $w^{(2)}$ schreiben:

$$a = y_1 \cdot w^{(1)} + y_2 \cdot w^{(2)}.$$

Zur Berechnung der Koeffizienten y_1, y_2 lösen wir das lineare Gleichungssystem

$$\begin{pmatrix} w_1^{(1)} & w_1^{(2)} \\ w_2^{(1)} & w_2^{(2)} \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} a_1 \\ a_2 \end{pmatrix}. \quad (3.23)$$

Beispiel 3.2.3 (Berechnung der Iterierten für einen speziellen Startwert).

Wir berechnen nun explizit die Werte von $v^{(k)}$ für das Beispiel $a = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ (ein junges Paar). Zur Darstellung des Vektors a als Linearkombination von $w^{(1)}$ und $w^{(2)}$ lösen wir (vgl. (3.23))

$$\begin{pmatrix} \frac{-1-\sqrt{5}}{2} & \frac{-1+\sqrt{5}}{2} \\ 1 & 1 \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}.$$

Die Lösung ist

$$y = \begin{pmatrix} \frac{-1}{\sqrt{5}} \\ \frac{1}{\sqrt{5}} \end{pmatrix},$$

also

$$a = \frac{-1}{\sqrt{5}}w^{(1)} + \frac{1}{\sqrt{5}}w^{(2)}.$$

Jetzt können wir den Zustand $v^{(k)}$ (Population im Monat k) berechnen:

$$\begin{aligned} v^{(k)} = A^k a &= A^k \left(\frac{-1}{\sqrt{5}}w^{(1)} + \frac{1}{\sqrt{5}}w^{(2)} \right) \\ &= \frac{-1}{\sqrt{5}}A^k w^{(1)} + \frac{1}{\sqrt{5}}A^k w^{(2)} \\ &= \frac{-1}{\sqrt{5}}\lambda_1^k w^{(1)} + \frac{1}{\sqrt{5}}\lambda_2^k w^{(2)} \\ &= \frac{1}{\sqrt{5}} \begin{pmatrix} \left(\frac{1-\sqrt{5}}{2} \right)^k \frac{1+\sqrt{5}}{2} + \left(\frac{1+\sqrt{5}}{2} \right)^k \frac{-1+\sqrt{5}}{2} \\ - \left(\frac{1-\sqrt{5}}{2} \right)^k + \left(\frac{1+\sqrt{5}}{2} \right)^k \end{pmatrix}. \end{aligned}$$

Man sieht jetzt leicht, dass z.B. die Zahl der alten Kaninchenpaare (und somit die Gesamtzahl der Paare) **(asymptotisch) exponentiell** wächst:

$$\begin{aligned} v_2^{(k)} &= \frac{1}{\sqrt{5}} \left(\frac{1-\sqrt{5}}{2} \right)^k + \frac{1}{\sqrt{5}} \left(\frac{1+\sqrt{5}}{2} \right)^k, \\ \lim_{k \rightarrow \infty} \frac{v_2^{(k)}}{\frac{1}{\sqrt{5}}\lambda_2^k} &= 1. \end{aligned} \tag{3.24}$$

Im Sinne von (3.24) gilt

$$v_2^{(k)} \approx \frac{1}{\sqrt{5}} \left(\frac{1+\sqrt{5}}{2} \right)^k.$$

Asymptotisch wächst die Zahl der alten Paare jeden Monat um den Faktor $\lambda_2 \approx 1,618 \dots < 2$. Man überlegt sich leicht, dass auch die Gesamtzahl der Kaninchenpaare asymptotisch jeden Monat mit diesem Faktor wächst. Die Gesamtzahl der Paare im Monat n ist nämlich gleich der Zahl der alten Paare im Monat $n+1$. Das Wachstum ist also auch für Modell 2 exponentiell, geschieht aber nicht so schnell wie in Modell 1.

3.2.1 Definition von Eigenwerten und Eigenvektoren

Wir liefern nun noch die exakten Definitionen bereits benutzter Begriffe nach.

Definition 3.2.4 (Eigenwert, Eigenvektor, Eigenraum).Sei $A \in \mathbb{R}^{n \times n}$.

1. $\lambda \in \mathbb{C}$ heißt **Eigenwert** von A , wenn es ein $v \in \mathbb{C}^n \setminus \{0\}$ gibt mit

$$Av = \lambda v.$$

2. Der Vektor v heißt dann **Eigenvektor** von A zum Eigenwert λ .
(Achtung: Der Nullvektor kann kein Eigenvektor sein!)

3. Der Untervektorraum

$$E_\lambda = \text{Kern}(A - \lambda \mathbb{I}_n) \subset \mathbb{C}^n$$

heißt **Eigenraum** zum Eigenwert λ . (Er besteht aus allen Eigenvektoren von A zum Eigenwert λ und dem Nullvektor.)

Bemerkung 3.2.5. Der Nullvektor ist zwar kein Eigenvektor, aber die Zahl 0 kann Eigenwert sein. 0 ist Eigenwert von $A \in \mathbb{R}^{n \times n}$ wenn A singulär ist, d.h. wenn $\text{Kern}(A) \neq \{0\}$. (Mit $\{0\}$ ist der Nullvektorraum gemeint.)

Satz 3.2.6 (Charakteristische Gleichung einer quadratischen Matrix).Die Eigenwerte von $A \in \mathbb{R}^{n \times n}$ sind die Lösungen der Gleichung (in der Variablen λ)

$$\det(A - \lambda \mathbb{I}_n) = 0.$$

Die Funktion $\det(A - \lambda \mathbb{I}_n)$ ist ein Polynom vom Grad n in λ , dessen Koeffizienten von den Einträgen (Koeffizienten) der Matrix A abhängen.

Definition 3.2.7 (Charakteristisches Polynom einer quadratischen Matrix).Das Polynom $\det(A - \lambda \mathbb{I}_n)$ heißt das **charakteristische Polynom** von $A \in \mathbb{R}^{n \times n}$.**Beispiel 3.2.8** (Charakteristisches Polynom einer (2×2) -Matrix).Sei $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \in \mathbb{R}^{2 \times 2}$. Dann gilt

$$\begin{aligned} \det(A - \lambda I_2) &= \det \begin{pmatrix} a - \lambda & b \\ c & d - \lambda \end{pmatrix} \\ &= (a - \lambda)(d - \lambda) - bc \\ &= \lambda^2 - \underbrace{(a + d)}_{\text{Spur } A} + \underbrace{ad - bc}_{\det A} \end{aligned}$$

Die Summe der Diagonalelemente von A ist die **Spur** von A und wird mit $\text{Spur } A$ bezeichnet.

Zur Definition und zur Berechnung von Determinanten von Matrizen in $\mathbb{R}^{n \times n}$ mit $n \geq 3$ verweisen wir auf Kapitel 3.1. Wir weisen nochmal ausdrücklich darauf hin, dass ein Eigenwert einer Matrix $A \in \mathbb{R}^{n \times n}$ auch eine nicht-reelle komplexe Zahl sein kann.

Beispiel 3.2.9 ((2×2) -Drehmatrix).

Wir betrachten die **Drehmatrix**

$$A = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix}.$$

Die Multiplikation $A \cdot v$ entspricht einer Drehung von $v \in \mathbb{R}^2$ um den Winkel α gegen den Uhrzeigersinn.

Wir betrachten nun speziell das Beispiel für den Drehwinkel $\alpha = \frac{\pi}{2}$. Es gilt $\sin \frac{\pi}{2} = 1$, $\cos \frac{\pi}{2} = 0$, also

$$A = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}, \text{ Spur } A = 0, \det A = 1.$$

Das charakteristische Polynom $P(\lambda) = \lambda^2 + 1$ hat die Nullstellen $\lambda_1 = i$ und $\lambda_2 = -i$.

Wir berechnen den Eigenraum E_{λ_1} . Dazu lösen wir:

$$\begin{aligned} \begin{pmatrix} -i & -1 \\ 1 & -i \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} &= \begin{pmatrix} 0 \\ 0 \end{pmatrix} \\ \Leftrightarrow \begin{pmatrix} -i & -1 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} &= \begin{pmatrix} 0 \\ 0 \end{pmatrix} \\ \Leftrightarrow -ix_1 - x_2 &= 0. \end{aligned}$$

Wir können also $x_2 \in \mathbb{C}$ beliebig wählen und $x_1 = ix_2$. So erhalten wir den Vektor $x_2 \cdot \begin{pmatrix} i \\ 1 \end{pmatrix}$. Jeder Eigenvektor zu λ_1 lässt sich so darstellen. Also

$$E_{\lambda_1} = \left\{ \begin{pmatrix} i \\ 1 \end{pmatrix} \cdot x_2 \mid x_2 \in \mathbb{C} \right\} \subset \mathbb{C}^2.$$

Analog dazu berechnen wir

$$E_{\lambda_2} = \left\{ \begin{pmatrix} -i \\ 1 \end{pmatrix} \cdot x_2 \mid x_2 \in \mathbb{C} \right\} \subset \mathbb{C}^2.$$

3.3 Basen und Koordinatensysteme

Die Begriffe des Eigenwerts und des Eigenvektors werden transparenter, wenn wir noch einmal einen Schritt zurück gehen und versuchen, die lineare Abbildung unabhängig von einer speziellen

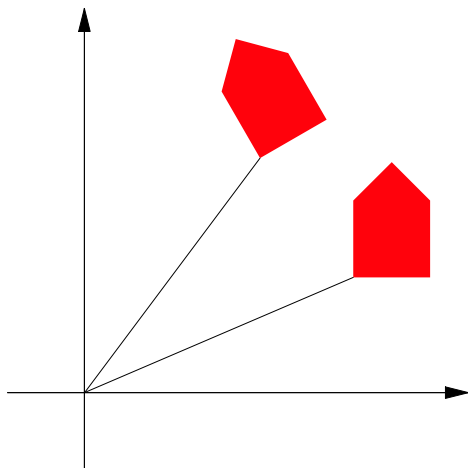


Abbildung 3.6: Eine Koordinatentransformation kann man sich entweder als Drehung (und evtl. Streckung) des Raumes vorstellen, die alle darin liegenden Objekte verändert...

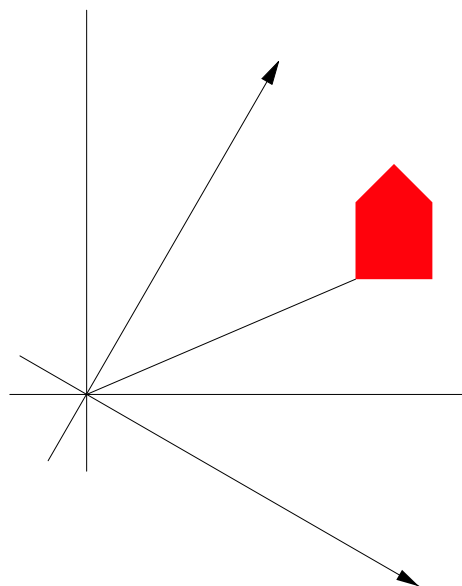


Abbildung 3.7: ...oder als Drehung des Koordinatensystems, wobei der Raum und alle darin liegenden Objekte an ihrem Platz verbleiben.

Basis zu betrachten. Wir behandeln nun also für einen Moment den $\mathbb{R}^n$ wie einen abstrakten Vektorraum.

In Kapitel 2.5.1 hatten wir bereits den Begriff des zu einer Basis gehörenden Koordinatensystems für Vektoren eingeführt, worauf wir nun zurückgreifen.

Seien $V = \mathbb{R}^n$ und $\mathcal{A} = (v_1, \dots, v_n)$ eine Basis mit Koordinatensystem

$$\begin{aligned}\phi_{\mathcal{A}} : \mathbb{R}^n &\rightarrow V \\ (x_1, \dots, x_n) &\mapsto x_1 v_1 + \dots + x_n v_n,\end{aligned}$$

sowie $\mathcal{B} = (w_1, \dots, w_n)$ eine zweite Basis von V mit Koordinationssystem

$$\begin{aligned}\phi_{\mathcal{B}} : \mathbb{R}^n &\rightarrow V \\ (y_1, \dots, y_n) &\mapsto y_1 w_1 + \dots + y_n w_n.\end{aligned}$$

Koordinatentransformation für Vektoren

Wie werden aus „alten“ Koordinaten $x = \phi_{\mathcal{A}}^{-1}(v)$ eines Vektors $v \in \mathbb{R}^n$ die „neuen“ Koordinaten $y = \phi_{\mathcal{B}}^{-1}(v)$ berechnet? Wie berechnet man also die Matrix, die der Abbildung $y = \phi_{\mathcal{B}}^{-1}(\phi_{\mathcal{A}}(x))$ entspricht? In Abbildungen 3.6 und 3.7 illustrieren wir eine Koordinatendrehung. In Abbildung 3.8 ist ein Vektor sowohl als Linearkombination von Standardbasisvektoren als auch von Vektoren einer anderen Basis dargestellt.

Zur Illustration betrachten wir das Beispiel aus Modell 2 zur Kaninchenpopulation.

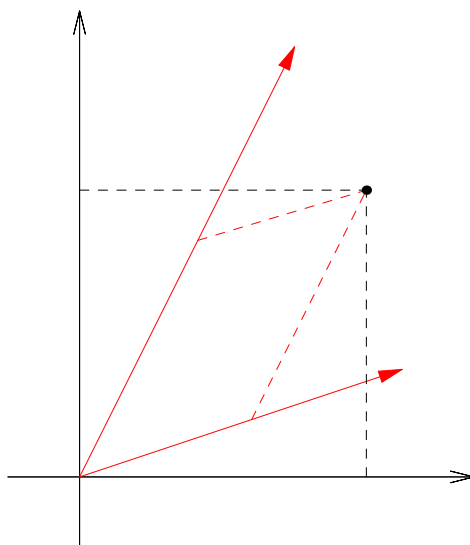


Abbildung 3.8: Darstellung eines Vektors in unterschiedlichen Basen

Beispiel 3.3.1 (Koordinatenwechsel für Modell 2).

Der Startvektor a aus Beispiel 3.2.3 hat bezüglich der Basis $\mathcal{A} = (e_1, e_2) = \left(\begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right)$ des $\mathbb{R}^2$ die Koordinaten $x = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$. Wir wählen nun als neue Basis $\mathcal{B} = (w^{(1)}, w^{(2)})$, wobei $w^{(i)}$ die Eigenvektoren aus (3.21) und (3.22) von $A = \begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix}$ zu den Eigenwerten $\lambda_1 = \frac{1-\sqrt{5}}{2}$ und $\lambda_2 = \frac{1+\sqrt{5}}{2}$, respektive, sind. Bezüglich der alten Basis $\mathcal{A}$ haben die neuen Basisvektoren folgende Darstellung:

$$\begin{aligned} w^{(1)} &= \frac{-1-\sqrt{5}}{2} \begin{pmatrix} 1 \\ 0 \end{pmatrix} + 1 \cdot \begin{pmatrix} 0 \\ 1 \end{pmatrix} \\ &= \begin{pmatrix} \frac{-1-\sqrt{5}}{2} \\ 1 \end{pmatrix}_{\mathcal{A}} \\ w^{(2)} &= \frac{-1+\sqrt{5}}{2} \begin{pmatrix} 1 \\ 0 \end{pmatrix} + 1 \cdot \begin{pmatrix} 0 \\ 1 \end{pmatrix} \\ &= \begin{pmatrix} \frac{-1+\sqrt{5}}{2} \\ 1 \end{pmatrix}_{\mathcal{A}}, \end{aligned}$$

wobei wir hier durch die Indizierung mit $\mathcal{A}$ explizit angeben, dass wir die Koordinatendarstellung bezüglich der Basis $\mathcal{A}$ meinen. Bezüglich der neuen Basis $\mathcal{B}$ hat a die Darstellung

$$\begin{aligned} a &= \begin{pmatrix} y_1 \\ y_2 \end{pmatrix}_{\mathcal{B}}, \text{ d.h.} \\ a &= y_1 \cdot w^{(1)} + y_2 w^{(2)}, \end{aligned} \tag{3.25}$$

wobei y_1 und y_2 noch zu bestimmen sind. Gleichung (3.25) für y_1, y_2 lässt sich in der Koordinatendarstellung bezüglich der alten Basis $\mathcal{A}$ wie folgt schreiben:

$$\begin{pmatrix} 1 \\ 0 \end{pmatrix} = y_1 \begin{pmatrix} \frac{-1-\sqrt{5}}{2} \\ 1 \end{pmatrix} + y_2 \begin{pmatrix} \frac{-1+\sqrt{5}}{2} \\ 1 \end{pmatrix}.$$

Wir müssen also folgendes lineare Gleichungssystem lösen:

$$\begin{pmatrix} \frac{-1-\sqrt{5}}{2} & \frac{-1+\sqrt{5}}{2} \\ 1 & 1 \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}.$$

Die Lösung ist:

$$\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} -\frac{1}{\sqrt{5}} \\ \frac{1}{\sqrt{5}} \end{pmatrix}.$$

Somit haben wir die Darstellung des Vektors a bezüglich zweier verschiedener Basen, $\mathcal{A}$ und $\mathcal{B}$, berechnet.

Allgemeiner linearer Koordinatenwechsel für Vektoren

Wir zeigen nun, wie man allgemein y aus x berechnet, wenn die Basen $\mathcal{A}$ und $\mathcal{B}$ gegeben sind. Seien also $\mathcal{A} = (v_1, \dots, v_n)$, $\mathcal{B} = (w_1, \dots, w_m)$. Der Koordinatenwechsel ist eine lineare Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^n$, ist also wie folgt durch eine Matrix S gegeben: Da $\mathcal{A}$ eine Basis des $\mathbb{R}^n$ ist, gibt es Koeffizienten s_{ij} ($1 \leq i, j \leq n$) mit

$$w_j = s_{1j}v_1 + s_{2j}v_2 + \dots + s_{nj}v_n.$$

Dadurch ist die Matrix $S = (s_{ij})_{1 \leq i, j \leq n}$ definiert.

Der Vektor $v \in V$ habe bezüglich $\mathcal{B}$ die Koordinaten $y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}$ und bezüglich $\mathcal{A}$ die Koordi-

naten $x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$. Dann gilt

$$\begin{aligned}
 v &= \Phi_{\mathcal{A}}(y) \\
 &= y_1 \omega_1 + y_2 \omega_2 + \cdots + y_n \omega_n \\
 &= y_1 (s_{11} v_1 + s_{21} v_2 + \cdots + s_{n1} v_n) \\
 &\quad + y_2 (s_{12} v_1 + s_{22} v_2 + \cdots + s_{n2} v_n) \\
 &\quad + \cdots \\
 &\quad + y_n (s_{1n} v_1 + s_{2n} v_2 + \cdots + s_{nn} v_n) \\
 &= (s_{11} y_1 + s_{12} y_2 + \cdots + s_{1n} y_n) \cdot v_1 \\
 &\quad (s_{21} y_1 + s_{22} y_2 + \cdots + s_{2n} y_n) \cdot v_2 \\
 &\quad + \cdots \\
 &\quad + (s_{n1} y_1 + s_{n2} y_2 + \cdots + s_{nn} y_n) \cdot v_n \\
 &\stackrel{!}{=} x_1 v_1 + \cdots + x_n v_n.
 \end{aligned}$$

Aus der letzten Gleichung erhalten wir durch Koeffizientenvergleich:

$$\begin{aligned}
 &\begin{cases} x_1 = s_{11} y_1 + \cdots + s_{1n} y_n \\ \vdots \\ x_n = s_{n1} y_1 + \cdots + s_{nn} y_n \end{cases} \\
 &\Leftrightarrow x = S y \\
 &\Leftrightarrow y = S^{-1} x.
 \end{aligned}$$

Wir fassen dieses Ergebnis im folgenden Satz zusammen.

Satz 3.3.2 (Linearer Koordinatenwechsel von Vektoren).

Seien V ein n -dim. reeller Vektorraum und $\mathcal{A} = (v_1, \dots, v_n)$ und $\mathcal{B} = (w_1, \dots, w_n)$ Basen von V mit Koordinatenabbildungen $\Phi_{\mathcal{A}}$ und $\Phi_{\mathcal{B}}$, respektive. Die Matrix $S = (s_{ij})_{1 \leq i, j \leq n} \in \mathbb{R}^{n \times n}$ sei durch

$$w_j = s_{1j} v_1 + \cdots + s_{nj} v_n \quad \forall 1 \leq j \leq n$$

bestimmt. In den Spalten von S stehen die Koeffizienten der Darstellung der (neuen) Basisvektoren w_i bezüglich der (alten) Basis $\mathcal{A}$. Ein Vektor $v \in V$ habe bezüglich $\mathcal{B}$ die Koordinaten

$$y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}, \text{ d.h. } v = \Phi_{\mathcal{B}}(y) = y_1 w_1 + \cdots + y_n w_n$$

und bezüglich $\mathcal{A}$ die Koordinaten

$$x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}, \text{ d.h. } v = \Phi_{\mathcal{A}}(x) = x_1 v_1 + \cdots + x_n v_n.$$

Dann ist der Koordinatenwechsel von y nach x durch

$$x = S y$$

gegeben und der von x nach y durch

$$y = S^{-1} x.$$

Definition 3.3.3. (Transformationsmatrix für linearen Koordinatenwechsel von Vektoren)

In der Situation von Satz 3.3.2 wird die Matrix

$$T_{\mathcal{A} \rightarrow \mathcal{B}} := S^{-1} \quad (3.26)$$

als **Transformationsmatrix** für den Basiswechsel von $\mathcal{A}$ nach $\mathcal{B}$ bezeichnet. Den Koordinatenvektor y eines Vektors bezüglich der neuen Basis $\mathcal{B}$ erhält man aus dessen Koordinatenvektor x bezüglich der alten Basis $\mathcal{A}$ durch Multiplikation mit $T_{\mathcal{A} \rightarrow \mathcal{B}}$ (s. Abbildung 3.9):

$$y = T_{\mathcal{A} \rightarrow \mathcal{B}} \cdot x. \quad (3.27)$$

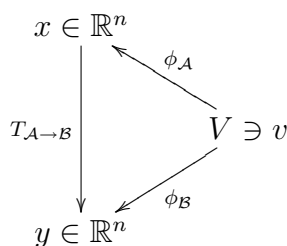


Abbildung 3.9: Kommutatives Diagramm zur Koordinatentransformation für Vektoren bei Basiswechsel von $\mathcal{A}$ zu $\mathcal{B}$

Beispiel 3.3.4 (Noch einmal: Koordinatenwechsel für Modell 2).

Vgl. Beispiel 3.3.1. Wir berechnen erneut die Koordinaten $y = \begin{pmatrix} y_1 \\ y_2 \end{pmatrix}$ des Startvektors a bezüglich der neuen Basis $\mathcal{B} = (w^{(1)}, w^{(2)})$. Diesmal gehen wir dabei ganz schematisch gemäß Satz 3.3.2 vor. Unsere Rechnung ist im Wesentlichen die gleiche wie in Beispiel 3.3.1, aber ihre Notation ist etwas kürzer und übersichtlicher.

Bezüglich $\mathcal{A}$ hat a die Koordinaten $\begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$. Es gilt:

$$w^{(1)} = \frac{-1 - \sqrt{5}}{2} \cdot \begin{pmatrix} 1 \\ 0 \end{pmatrix} + 1 \cdot \begin{pmatrix} 0 \\ 1 \end{pmatrix} \quad (\text{Diese Gleichung liefert die 1. Spalte von } S),$$

$$w^{(2)} = \frac{-1 + \sqrt{5}}{2} \cdot \begin{pmatrix} 1 \\ 0 \end{pmatrix} + 1 \cdot \begin{pmatrix} 0 \\ 1 \end{pmatrix} \quad (\text{Diese Gleichung liefert die 2. Spalte von } S).$$

Also

$$S = \begin{pmatrix} \frac{-1-\sqrt{5}}{2} & \frac{-1+\sqrt{5}}{2} \\ 1 & 1 \end{pmatrix}$$

$$S^{-1} = \frac{1}{\sqrt{5}} \begin{pmatrix} -1 & \frac{-1+\sqrt{5}}{2} \\ 1 & \frac{1+\sqrt{5}}{2} \end{pmatrix},$$

und somit

$$\begin{aligned} y &= S^{-1}x \\ &= +\frac{1}{\sqrt{5}} \begin{pmatrix} -1 & \frac{-1+\sqrt{5}}{2} \\ 1 & \frac{1+\sqrt{5}}{2} \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix} \\ &= \begin{pmatrix} \frac{-1}{\sqrt{5}} \\ \frac{1}{\sqrt{5}} \end{pmatrix}. \end{aligned}$$

Dies stimmt mit dem Ergebnis aus Beispiel 3.3.1 überein.

3.3.1 Koordinatentransformation für lineare Abbildungen

Vektoren $v \in V$ werden durch Koordinaten (n -Tupel, Elemente des $\mathbb{R}^n$) dargestellt, die durch die Wahl einer Basis $\mathcal{A}_1$ eindeutig definiert sind (siehe Kapitel 2.5.1). Und lineare Abbildungen $f : V \rightarrow W$ werden durch Matrizen dargestellt, die durch die Wahl von Basen $\mathcal{A}_1$ von V und $\mathcal{B}_1$ von W eindeutig definiert sind (siehe Satz 2.6.2). Wir wissen bereits, wie die Koordinaten von $v \in V$ bei Basiswechsel von $\mathcal{A}_1$ zu $\mathcal{A}_2$ und von $w \in W$ bei Basiswechsel von $\mathcal{B}_1$ zu $\mathcal{B}_2$ transformiert werden.

Im folgenden Satz zeigen wir, wie man die Darstellung von f bezüglich der neuen Basen aus der Darstellung von f bezüglich der alten Basen berechnet.

Satz 3.3.5. (Koordinatentransformation für lineare Abbildungen)

Sei $f : V \rightarrow W$ eine lineare Abbildung zwischen reellen Vektorräumen. Die Koordinatentransformation für Vektoren in V bei Basiswechsel von $\mathcal{A}_1$ nach $\mathcal{A}_2$ seien durch die Transformationsmatrix $T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2}$ beschrieben (vgl. Definition 3.3.3), und die Koordinatentransformation für Vektoren in W bei Basiswechsel von $\mathcal{B}_1$ nach $\mathcal{B}_2$ durch die Transformationsmatrix $T_{\mathcal{B}_1 \rightarrow \mathcal{B}_2}$. Sei des Weiteren f bezüglich der alten Basen $\mathcal{A}_1$ und $\mathcal{B}_1$ durch die Matrix A dargestellt.

Dann wird f bezüglich der neuen Basen $\mathcal{A}_2$ und $\mathcal{B}_2$ durch die Matrix

$$T_{\mathcal{B}_1 \rightarrow \mathcal{B}_2} \cdot A \cdot T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2}^{-1} \tag{3.28}$$

dargestellt.

Beweis: Sei $\dim V = n$ und $\dim W = m$. Gleichung (3.28) liest man einfach aus dem *kommutativen Diagramm* in Abbildung 3.10 ab: Man gelangt von links unten nach rechts unten auf zwei verschiedenen Wegen, einmal direkt entlang dem horizontalen Pfeil- dieser entspricht der

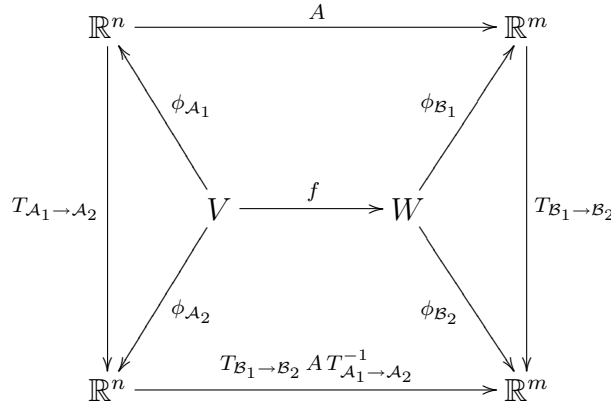


Abbildung 3.10: Kommutatives Diagramm zur Koordinatentransformation für lineare Abbildungen bei Basiswechsel von $\mathcal{A}$ zu $\mathcal{B}$

Matrix, welche f bezüglich der neuen Koordinaten darstellt- und einmal indirekt: erst nach oben (entspricht der Inversen von $T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2}$), dann horizontal nach rechts (entspricht der Matrix A , die f bezüglich der alten Koordinaten darstellt) und dann nach unten (entspricht der Matrix $T_{\mathcal{B}_1 \rightarrow \mathcal{B}_2}$). Da das Diagramm kommutativ ist und beide Wege denselben Anfangspunkt und denselben Endpunkt haben, entsprechen sie den gleichen Matrizen, wobei der zweite Weg dem Produkt der drei genannten Matrizen entspricht. Es folgt also Formel (3.28).

Beweis (2. Version): Wir geben noch einen alternativen Beweis mit Formeln an, der aber im Wesentlichen völlig analog verläuft: Seien $v \in V$ und

$$f(v) = w \in W \quad (3.29)$$

Wir betrachten zunächst die Darstellung von Gleichung (3.29) in Koordinaten bezüglich der alten Basen. Bezüglich $\mathcal{A}_1$ werde v durch den Koordinatenvektor $x^{(1)} \in \mathbb{R}^n$, bezüglich $\mathcal{B}_1$ werde w durch den Koordinatenvektor $y^{(1)} \in \mathbb{R}^m$, und die lineare Abbildung f werde durch $A \in \mathbb{R}^{m \times n}$ dargestellt. Also ist Gleichung (3.29) äquivalent zu Gleichung (3.30).

$$Ax^{(1)} = y^{(1)} \quad (3.30)$$

$$\Leftrightarrow T_{\mathcal{B}_1 \rightarrow \mathcal{B}_2} A T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2}^{-1} T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2} x^{(1)} = T_{\mathcal{B}_1 \rightarrow \mathcal{B}_2} y^{(1)} \quad (3.31)$$

$$\Leftrightarrow T_{\mathcal{B}_1 \rightarrow \mathcal{B}_2} A T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2}^{-1} x^{(2)} = y^{(2)}. \quad (3.32)$$

Im Schritt von (3.30) nach (3.31) haben wir beide Seiten von links mit der regulären Matrix $T_{\mathcal{B}_1 \rightarrow \mathcal{B}_2}$ multipliziert und auf der linken Seite zwischen A und $x^{(1)}$ die identische Matrix $T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2}^{-1} T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2}$ eingefügt. Für den Schritt von (3.31) nach (3.32) haben wir den Koordinatenvektor von v bezüglich der neuen Basis $\mathcal{A}_2$ mit $x^{(2)}$ und den Koordinatenvektor von $f(v)$ bezüglich der neuen Basis $\mathcal{B}_2$ mit $y^{(2)}$ bezeichnet und die Identitäten $x^{(2)} = T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2} x^{(1)}$ und $y^{(2)} = T_{\mathcal{B}_1 \rightarrow \mathcal{B}_2} y^{(1)}$ verwendet. Damit ist offensichtlich Gleichung (3.32) die Darstellung von Gleichung (3.29) im neuen Koordinatensystem und die darstellende Matrix ist die aus Formel (3.28). $\square$

Beispiel 3.3.6 (Transformation der Matrix zu Modell 2).

Wir betrachten wieder das Beispiel von Modell 2.

$$\begin{aligned} A &= \begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix}, \\ \mathcal{A}_1 &= \mathcal{B}_1 = \left(\begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right), \\ \mathcal{A}_2 &= \mathcal{B}_2 = (w^{(1)}, w^{(2)}). \end{aligned}$$

Der Koordinatenwechsel für Vektoren, von Basis $\mathcal{A}_1$ zu $\mathcal{A}_2$ ist durch die Matrix $T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2}$ gegeben:

$$\begin{aligned} T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2}^{-1} &= S = \begin{pmatrix} \frac{-1-\sqrt{5}}{2} & \frac{-1+\sqrt{5}}{2} \\ 1 & 1 \end{pmatrix}, \\ T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2} &= S^{-1} = \frac{1}{\sqrt{5}} \begin{pmatrix} -1 & \frac{-1+\sqrt{5}}{2} \\ 1 & \frac{-1-\sqrt{5}}{2} \end{pmatrix}. \end{aligned}$$

Wir berechnen die darstellende Matrix bezüglich der neuen Basis $\mathcal{A}_2 = \mathcal{B}_2$:

$$T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2} \cdot A \cdot T_{\mathcal{A}_1 \rightarrow \mathcal{A}_2}^{-1} = \begin{pmatrix} \frac{1-\sqrt{5}}{2} & 0 \\ 0 & \frac{1+\sqrt{5}}{2} \end{pmatrix}.$$

3.3.2 Ähnlichkeit von Matrizen

An einigen Beispielen von linearen dynamischen Systemen wie z.B. Kaninchenpopulationen, Mischen von Lösungen (s. Hausaufgaben), die hier durch lineare Abbildungen $f : V \rightarrow V$ gegeben sind, sehen wir, dass das Langzeitverhalten (Verhalten $f^n v$ für $v \in V$ und „grosse“ $n \in \mathbb{N}$) solcher Systeme durch die Eigenwerte der darstellenden Matrix charakterisiert wird. Eine solche Matrix hängt aber von der Wahl des Koordinatensystems (der Basis) ab. Für die Basis $\mathcal{A}$ von V werde f durch die Matrix $A \in \mathbb{R}^{n \times n}$ beschrieben. Bei Wahl einer anderen Basis $\mathcal{B}$ werde f durch die Matrix $B \in \mathbb{R}^{n \times n}$ dargestellt, wobei $B = TAT^{-1}$ ist und T den Koordinatenwechsel beschreibt.

Definition 3.3.7 (Ähnlichkeit von Matrizen).

Seien $A, B \in \mathbb{R}^{n \times n}$. A und B heißen einander **ähnlich**, wenn es eine reguläre Matrix $T \in \mathbb{R}^{n \times n}$ gibt mit $B = TAT^{-1}$.

Satz 3.3.8 (Ähnliche Matrizen haben das gleiche charakteristische Polynom).

Seien $A, B \in \mathbb{R}^{n \times n}$ ähnlich. Dann haben A und B das gleiche charakteristische Polynom und somit insbesondere auch die gleichen Eigenwerte.

Beweis: Sei $B = TAT^{-1}$. Dann gilt wegen $\det(T^{-1}) = (\det(t))^{-1}$:

$$\begin{aligned} \det(B - \lambda I) &= \det(TAT^{-1} - T \cdot \lambda I \cdot T^{-1}) \\ &= \det(T(A - \lambda I)T^{-1}) \\ &= \det(T) \cdot \det(A - \lambda I) \cdot \det(T^{-1}) \\ &= \det(A - \lambda I). \end{aligned}$$

□

Wir können also von den Eigenwerten des Endomorphismus bzw. des linearen Systems sprechen, da diese nicht von der speziellen Wahl der Koordinaten abhängen. Die hier vorgestellte Theorie wird uns insbesondere im Kapitel über Dynamische Systeme wiederbegegnen.

3.3.3 Diagonalisierbarkeit

Allgemein nennt man jede Matrix A , für die man eine Basis finden kann, bezüglich der sie durch eine Diagonalmatrix dargestellt wird, **diagonalisierbar**.

Definition 3.3.9 (Diagonalisierbarkeit).

Eine quadratische Matrix $A \in \mathbb{R}^{n \times n}$ heißt **diagonalisierbar**, wenn es eine Basis $(v_1, \dots, v_n)$ des $\mathbb{R}^n$ gibt, die nur aus Eigenvektoren der Matrix A besteht. Schreibt man die Eigenvektoren als Spalten in eine Matrix $S := (v_1 | \dots | v_n)$, so hat die Matrix $D = S^{-1}AS$ Diagonalgestalt (A und die Diagonalmatrix D sind also **ähnlich** zueinander).

Man kann die Relation zwischen A und D natürlich auch ausnutzen, um A darzustellen als

$$A = SDS^{-1} = \left(v_1 \mid \dots \mid v_n \right) \begin{pmatrix} \lambda_1 & & \\ & \ddots & \\ & & \lambda_n \end{pmatrix} \left(v_1 \mid \dots \mid v_n \right)^{-1}$$

und die Interpretation des Ausdrucks $A = SDS^{-1}$ ist die folgende: will man für einen beliebigen Vektor $v \in \mathbb{R}^n$ den Ausdruck Av berechnen, so kann man zunächst die Koordinaten von v in der neuen Basis (die durch die Spaltenvektoren von S gegeben ist), d.h. den Koordinatenvektor $S^{-1}v$ berechnen. In dieser Basis hat der Operator A Diagonalgestalt und wird durch die Diagonalmatrix D ausgedrückt, d.h. $DS^{-1}v$ ergibt bereits die Koordinaten von Av in der Basis S . Um jetzt das Ergebnis in der ursprünglichen (kanonischen) Basis zu erhalten, müssen wir nur noch den bereits berechneten Koordinatenvektor mit der Matrix S multiplizieren: so erhalten wir $SDS^{-1}v = Av$.

Bemerkung 3.3.10 (Vorteile von Diagonalmatrizen).

In Beispiel 3.3.6 haben wir durch den Wechsel zu einer Basis aus Eigenvektoren von A erreicht, dass die lineare Abbildung bezüglich der neuen Basis durch eine Diagonalmatrix

$$D = TAT^{-1} = \begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix}$$

dargestellt wird, deren Diagonalelemente gerade die Eigenwerte von A (und von D) sind. Mit Hilfe dieser können wir leicht Potenzen A^n von A und somit von $A^n x$ ausrechnen.

Es gilt:

$$\begin{aligned} D &= TAT^{-1} \\ \Leftrightarrow A &= T^{-1}DT. \end{aligned}$$

Also

$$\begin{aligned}
 A^n &= A \cdot A \cdot \dots \cdot A \\
 &= T^{-1} D \underbrace{TT^{-1}}_{\mathbb{I}} DT \dots T^{-1} DT \\
 &= T^{-1} D^n T \\
 &= T^{-1} \cdot \begin{pmatrix} \lambda_1^n & 0 \\ 0 & \lambda_2^n \end{pmatrix} \cdot T.
 \end{aligned}$$

Ebenso

$$A^n x = T^{-1} D^n T x.$$

Beispiel 3.3.11 (Diagonalisierung einer symmetrischen (2×2) -Matrix).

Zur Einübung der Transformation von Matrizen bei Basiswechsel diagonalisieren wir die symmetrische Matrix $A = \begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix}$.

Dazu berechnen wir die Eigenwerte und eine Basis von Eigenvektoren von A :

Das charakteristische Polynom von A ist

$$\begin{aligned}
 P(\lambda) &= \det(A - \lambda I_2) \\
 &= \det \begin{pmatrix} 1 - \lambda & 2 \\ 2 & 1 - \lambda \end{pmatrix} \\
 &= (1 - \lambda)^2 - 4 \\
 &= (\lambda + 1)(\lambda - 3).
 \end{aligned}$$

Die Eigenwerte von A sind die Nullstellen von P , also $\lambda_1 = -1$ und $\lambda_2 = 3$.

Eigenraum zu $\lambda_1 = -1$: Zu lösen ist das lineare Gleichungssystem $(A - \lambda_1 I_2)x = 0$ in den Variablen $x_1, x_2 \in \mathbb{C}$, also

$$\begin{aligned}
 \begin{pmatrix} 2 & 2 \\ 2 & 2 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} &= \begin{pmatrix} 0 \\ 0 \end{pmatrix} \\
 \Leftrightarrow \begin{pmatrix} 2 & 2 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} &= \begin{pmatrix} 0 \\ 0 \end{pmatrix} \\
 \Leftrightarrow 2x_1 + 2x_2 &= 0.
 \end{aligned}$$

Man kann $x_1 \in \mathbb{C}$ beliebig wählen und dann $x_2 = -x_1$. Also

$$E_{\lambda_1} = \left\{ x_1 \begin{pmatrix} -1 \\ 1 \end{pmatrix} \mid x_1 \in \mathbb{C} \right\}.$$

Eigenraum zu $\lambda_2 = 3$: Zu lösen ist

$$\begin{aligned}
 \begin{pmatrix} 1 - 3 & 2 \\ 2 & 1 - 3 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} &= \begin{pmatrix} 0 \\ 0 \end{pmatrix} \\
 \Leftrightarrow 2x_1 - 2x_2 &= 0.
 \end{aligned}$$

$$E_{\lambda_2} = \left\{ x_1 \begin{pmatrix} 1 \\ 1 \end{pmatrix} \mid x_1 \in \mathbb{C} \right\}.$$

Wir wählen nun aus jedem Eigenraum einen Vektor und erhalten mit $\left(\begin{pmatrix} -1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right)$ eine Orthogonalbasis des $\mathbb{R}^2$. Die Spalten der Matrix $S = T^{-1}$ sind die Koordinatenvektoren (bezüglich der alten Basis) der neuen Basisvektoren, also

$$T^{-1} = S = \begin{pmatrix} -1 & 1 \\ 1 & 1 \end{pmatrix}.$$

Wir erhalten T durch Invertierung von S :

$$T = \begin{pmatrix} -\frac{1}{2} & \frac{1}{2} \\ \frac{1}{2} & \frac{1}{2} \end{pmatrix}.$$

Somit hat der bezüglich der Standardbasis durch A dargestellte Endomorphismus bzgl. der neuen (orthogonalen) Basis die Darstellung

$$\begin{aligned} TAT^{-1} &= \begin{pmatrix} -\frac{1}{2} & \frac{1}{2} \\ \frac{1}{2} & \frac{1}{2} \end{pmatrix} \underbrace{\begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix} \begin{pmatrix} -1 & 1 \\ 1 & 1 \end{pmatrix}} \\ &= \begin{pmatrix} -\frac{1}{2} & \frac{1}{2} \\ \frac{1}{2} & \frac{1}{2} \end{pmatrix} \begin{pmatrix} 1 & 3 \\ -1 & 3 \end{pmatrix} \\ &= \begin{pmatrix} -1 & 0 \\ 0 & 3 \end{pmatrix}. \end{aligned}$$

Die Diagonalelemente dieser Matrix sind natürlich die Eigenwerte von A .

Wir bemerken noch, dass in den Spalten der Matrix T die Koordinaten (bezüglich der neuen Basis) der (alten) Standardbasisvektoren stehen.

Kapitel 4

Analysis über $\mathbb{R}$

Schon im alten Griechenland war einigen Mathematikern aufgefallen, dass die Menge der rationalen Zahlen (also die Menge der Brüche $\frac{p}{q}$ mit $p, q \in \mathbb{Z}$), die wir heute $\mathbb{Q}$ nennen, „Lücken“ hat. Will man die Länge x der Diagonalen eines Quadrates mit der Seitenlänge 1 berechnen, so gelangt man mit Hilfe des Satzes von Pythagoras zur Gleichung $1^2 + 1^2 = x^2$. Man kann aber zeigen, dass die Gleichung $x^2 = 2$ keine positive rationale Lösung hat. Wir können aber die 2 durch Quadrate von rationalen Zahlen beliebig eng einschachteln, z.B. durch bestapproximierende Dezimalbrüche vorgegebener Länge:

$$1^2 < 1.4^2 < 1.41^2 < 1.414^2 < \dots < 2 < \dots < 1.415^2 < 1.42^2 < 1.5^2 < 2^2. \quad (4.1)$$

Und daraus erhalten wir eine aufsteigende und eine absteigende Folge von rationalen Zahlen:

$$\begin{array}{ccccccc} 1 & < & 1.4 & < & 1.41 & < & 1.414 & < & \dots \\ 2 & > & 1.5 & > & 1.42 & > & 1.415 & > & \dots \end{array}$$

Obwohl sämtliche Glieder der ersten Folge kleiner sind als alle Glieder der zweiten Folge, die beide Folgen also separiert sind, gibt es keine rationale Zahl, die zwischen ihnen liegt. Durch das „Stopfen“ solcher Lücken gelangt man von den rationalen Zahlen zur Menge $\mathbb{R}$ der reellen Zahlen, den für den Anwender vielleicht wichtigsten Zahlen der Mathematik, mit denen wir üblicherweise rechnen und es in dieser Vorlesung ja bereits ausgiebig getan haben. In Kapitel 5 werden wir noch einen weiteren wichtigen Zahltyp behandeln, die *komplexen Zahlen*.

Hier in diesem Kapitel beschäftigen wir uns mit *Folgen* und *Reihen* reeller Zahlen, *Stetigkeit* und *Differentiation* von Funktionen sowie den für die Praxis äußerst wichtigen *Taylorreihen*.

4.1 Folgen und Konvergenz

Wir betrachten nun also Folgen von reellen Zahlen:

Definition 4.1.1 (Folge).

Eine **Folge** a mit Werten in $\mathbb{R}$ ist eine Abbildung

$$\begin{aligned} a : \mathbb{N} &\longrightarrow \mathbb{R}, \\ n &\longmapsto a(n). \end{aligned}$$

Wir schreiben auch a_n statt $a(n)$ für das Folgenglied mit Index n , und die gesamte Folge bezeichnen wir auch mit $(a_n)_{n \in \mathbb{N}}$ oder $(a_n)_{n \geq 0}$ oder, je nach Indexmenge, z.B. auch $(a_n)_{n \geq n_0}$. Zuweilen indizieren wir Folgenglieder auch mit einem hochgesetzten Index, also z.B. $(x^{(n)})_{n \in \mathbb{N}}$. Dabei setzen wir den Index n in Klammern, um Verwechslung mit x^n („ x hoch n “) zu vermeiden.

Definition 4.1.2 (Nullfolge).

Eine Folge $(a_n)_{n \in \mathbb{N}}$ heißt **Nullfolge**, wenn es für alle $\epsilon > 0$ ein $n_0 \in \mathbb{N}$ gibt, so dass für alle $n \geq n_0$ gilt:

$$|a_n| \leq \epsilon.$$

In Quantorenschreibweise lautet die Bedingung:

$$\forall \epsilon > 0 \quad \exists n_0 \quad \forall n \geq n_0 \quad |a_n| \leq \epsilon. \quad (4.2)$$

Wir sagen auch, die Folge $(a_n)_{n \in \mathbb{N}}$ **konvergiert gegen 0** oder die Folge hat den **Grenzwert 0**, und schreiben

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

Bemerkung 4.1.3. Wenn $(a_n)_{n \in \mathbb{N}}$ eine Nullfolge ist, muss es aber nicht unbedingt ein n mit $a_n = 0$ geben, wie das folgende Beispiel 4.1.4 zeigt.

Beispiel 4.1.4. Sei $a_n = \frac{1}{n}$. Dann ist $(a_n)_{n \geq 1}$ eine Nullfolge.

Beweis: Sei $\epsilon > 0$ gegeben. Wann ist die gewünschte Ungleichung

$$\frac{1}{n} \leq \epsilon \quad (4.3)$$

erfüllt? Bedingung (4.3) ist äquivalent zu

$$\frac{1}{\epsilon} \leq n.$$

Wir wählen ein n_0 mit $\frac{1}{\epsilon} \leq n_0$. Dann gilt für alle $n \geq n_0$:

$$\left| \frac{1}{n} \right| \leq \left| \frac{1}{n_0} \right| \leq \epsilon.$$

Da wir also für ein beliebiges ϵ ein (von ϵ abhängiges) n_0 finden können, welches (4.2) erfüllt, ist $(a_n)_{n \geq 1}$ eine Nullfolge. $\square$

Beispiel 4.1.5. Sei $a_n = \frac{1}{2^n}$. Die Folge $(\frac{1}{2^n})_{n \in \mathbb{N}}$ konvergiert gegen 0.

Beweis: (Gleiche Beweisführung wie bei Beispiel 4.1.4): Sei $\epsilon > 0$ gegeben:

Die Bedingung für die Folgeindizes n ist

$$\begin{aligned} \frac{1}{2^n} &\leq \epsilon \\ \Leftrightarrow \frac{1}{\epsilon} &\leq 2^n \end{aligned}$$

Zunächst überlegen wir uns, dass $2^n \geq n$ für $n \geq 0$. Dies folgt aus der Bernoulli-Ungleichung mit $a = 1$. Nach Beispiel 4.1.4 gibt es ein $n_0 \geq 2$, so dass für alle $n \geq n_0$ die Abschätzung

$$\frac{1}{\epsilon} \leq n$$

gilt, also wegen $2^n \geq n$ erst recht

$$\frac{1}{\epsilon} \leq 2^n.$$

□

Bemerkung 4.1.6 (Majorante).

Im Beweis haben wir eine **Majorante** $(a'_n)_{n \geq 1} = (\frac{1}{n})_{n \geq 1}$ von $(a_n)_{n \geq 1} = (\frac{1}{2^n})_{n \geq 1}$ verwendet, d.h. die zu untersuchende Folge wird von zwei Nullfolgen eingeschachtelt, der konstanten Nullfolge und der Majorante:

$$0 \leq a_n \leq a'_n.$$

Definition 4.1.7 (Konvergenz und Grenzwert einer Folge).

Eine Folge $(a_n)_{n \in \mathbb{N}}$ **konvergiert** gegen g , wenn gilt:

$$\forall \epsilon > 0 \quad \exists n_0 \in \mathbb{N} \quad \forall n \geq n_0 \quad |a_n - g| \leq \epsilon.$$

Wir bezeichnen g als **Grenzwert** der Folge und schreiben

$$\lim_{n \rightarrow \infty} a_n = g.$$

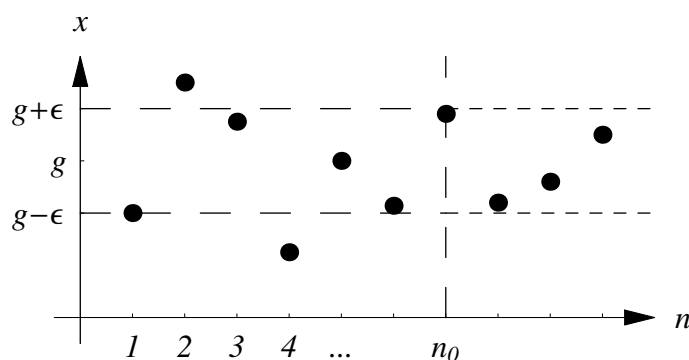


Abbildung 4.1: Wenn n_0 groß genug gewählt wird, liegen für alle $n \geq n_0$ die Folgenglieder a_n zwischen $g - \epsilon$ und $g + \epsilon$ für beliebiges $\epsilon > 0$.

Bemerkung 4.1.8. Es folgt sofort aus den Definitionen 4.1.2 und 4.1.7, dass eine Folge (a_n) genau dann gegen g konvergiert, wenn $(a_n - g)_{n \in \mathbb{N}}$ eine Nullfolge ist.

Satz 4.1.9. (Rechenregeln für Grenzwerte konvergenter Folgen)

Seien $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ konvergente Folgen mit $\lim_{n \rightarrow \infty} a_n = a$ und $\lim_{n \rightarrow \infty} b_n = b$ und $\lambda \in \mathbb{R}$. Dann gilt:

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

2.

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

3. speziell:

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

$$\lim_{n \rightarrow \infty} (a_n - b_n) = a - b,$$

$$\lim_{n \rightarrow \infty} (\lambda a_n) = \lambda a.$$

4.

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

5. Falls $a \neq 0$, dann ist für ein hinreichend großes n_0 die Folge $(\frac{1}{a_n})_{n \geq n_0}$ definiert und

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

6. Wenn die Voraussetzung von (5.) erfüllt ist und $\lim b_n = b$, dann ist

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

7. Ist $(c_n)_{n \in \mathbb{N}}$ eine beschränkte Folge und $\lim_{n \rightarrow \infty} b_n = 0$, dann

$$\lim_{n \rightarrow \infty} c_n \cdot b_n = 0.$$

Beweis: (nur exemplarisch):

(zu 2.) Sei $\epsilon > 0$ gegeben. Es gibt es ein n_0 und ein n_1 mit

$$\begin{array}{lll} |a_n - a| & \leq & \frac{\epsilon}{2|\lambda|} \quad \forall n > n_0 \\ \text{und } |b_n - b| & \leq & \frac{\epsilon}{2} \quad \forall n > n_1, \end{array}$$

und für alle $n \geq \max\{n_0, n_1\} =: n_3$ gilt

$$\begin{aligned} |(\lambda a_n + b_n) - (\lambda a + b)| &= |\lambda(a_n - a) + (b_n - b)| \\ &\leq \underbrace{|\lambda| \cdot |a_n - a|}_{\leq \frac{\epsilon}{2}, \text{ da } n \geq n_0} + \underbrace{|b_n - b|}_{\leq \frac{\epsilon}{2}, \text{ da } n \geq n_1} \\ &\leq \epsilon. \end{aligned}$$

(zu 3.) Die Aussagen sind Spezialfälle von (2.)

(zu 4.) Da die Folge $(b_n)_{n \in \mathbb{N}}$ konvergent und $(|b_n|)_{n \in \mathbb{N}}$ nach (1.) durch eine Konstante B beschränkt ist, gilt

$$\begin{aligned} |(a_n \cdot b_n) - ab| &= |a_n b_n - ab_n + ab_n - ab| \\ &\leq \underbrace{|b_n|}_{\leq B} \cdot |a_n - a| + |a| \cdot |b_n - b|. \end{aligned} \quad (4.4)$$

Wähle n_0 so, dass für alle $n \geq n_0$ die beiden folgenden Abschätzungen erfüllt sind:

$$\begin{aligned} |a_n - a| &\leq \frac{\epsilon}{2B}, \\ |b_n - b| &\leq \frac{\epsilon}{2 \cdot \max\{|a|, 1\}}. \end{aligned}$$

Dann folgt

$$\begin{aligned} \underbrace{|b_n|}_{\leq B} \cdot |a_n - a| + |a| \cdot |b_n - b| &\leq \frac{\epsilon}{2} + \frac{\epsilon}{2} \\ &= \epsilon. \end{aligned}$$

□

Definition 4.1.10 (monotone Folge).

Eine Folge $(a_n)_{n \geq n_0}$ heißt

1. **monoton steigend**, wenn für alle $n \geq n_0$ gilt: $a_n \leq a_{n+1}$.
2. **streng monoton steigend**, wenn für alle $n \geq n_0$ gilt: $a_n < a_{n+1}$.
3. **monoton fallend**, wenn für alle $n \geq n_0$ gilt: $a_n \geq a_{n+1}$.
4. **streng monoton fallend**, wenn für alle $n \geq n_0$ gilt: $a_n > a_{n+1}$.

Definition 4.1.11 (Cauchy-Folge).

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

$$\forall \epsilon > 0 \quad \exists n_0 \quad \forall n, m \geq n_0 \quad |a_n - a_m| \leq \epsilon.$$

Satz 4.1.12. (Konvergenz von Cauchy-Folgen und monotonen, beschränkten Folgen)

1. Jede Cauchy-Folge mit Werten in $\mathbb{R}$ oder $\mathbb{C}$ ist konvergent. Und jede konvergente Folge mit Werten in $\mathbb{R}$ oder $\mathbb{C}$ ist eine Cauchyfolge.
2. Jede reelle nach oben beschränkte, monoton steigende Folge ist konvergent.
Jede reelle nach unten beschränkte, monoton fallende Folge ist konvergent.

Bemerkung: Die Kriterien aus Satz 4.1.12 können sehr nützlich zum Nachweis der Konvergenz sein, wenn der Grenzwert nicht bekannt ist.

Beispiel 4.1.13. (Eulersche Zahl als Grenzwert einer Folge)

Betrachte die durch $a_n := (1 + \frac{1}{n})^n$ für $n \geq 1$ definierte Folge.

1. $(a_n)_{n \geq 1}$ ist monoton steigend.

Beweis:

$$\begin{aligned} \frac{a_n}{a_{n-1}} &= \left(\frac{n+1}{n}\right)^n \cdot \left(\frac{n-1}{n}\right)^{n-1} = \left(\frac{n^2-1}{n^2}\right)^n \cdot \frac{n}{n-1} \\ &= \left(1 - \frac{1}{n^2}\right)^n \cdot \frac{n}{n-1} \\ &\geq \left(1 - \frac{1}{n}\right) \cdot \frac{n}{n-1} = 1, \end{aligned}$$

wobei wir die Bernoulli-Ungleichung (s. Satz 1.4.5) verwendet haben.

2. Ebenso zeigt man, dass für $b_n = (1 + \frac{1}{n})^{n+1}$ die Abschätzung

$$0 \leq a_n \leq b_n$$

gilt und $(b_n)_{n \in \mathbb{N}}$ eine monoton fallende Folge ist, also insbesondere

$$a_n \leq b_1 = 4.$$

Also ist $(a_n)_{n \in \mathbb{N}}$ monoton steigend und nach oben beschränkt. Nach Satz 4.1.12.2 hat $(a_n)_n$ einen Grenzwert.

Dieser Grenzwert heißt **Eulersche Zahl** und wird mit e bezeichnet. Diese Zahl ist nicht rational, d.h. ihr Dezimalbruch ist nicht periodisch.

$$\boxed{\lim \left(1 + \frac{1}{n}\right)^n = e = 2.7182818285 \dots \quad (\text{Eulersche Zahl})} \quad (4.5)$$

Definition 4.1.14 (Divergenz einer Folge).

1. Eine Folge heißt **divergent**, wenn sie nicht konvergiert.
2. Eine reellwertige Folge $(a_n)_{n \in \mathbb{N}}$ **geht gegen** ∞ , wenn

$$\forall M > 0 \exists n_0 \in \mathbb{N} \forall n > n_0 \quad a_n > M.$$

Wir schreiben dann

$$\lim_{n \rightarrow \infty} a_n = \infty. \quad (4.6)$$

Analog dazu definieren wir, wann eine Folge **gegen** $-\infty$ **geht**.

Bemerkung 4.1.15.

1. Insbesondere sind Folgen divergent, die gegen ∞ oder gegen $-\infty$ gehen. Die Umkehrung gilt nicht. Es gibt z.B. beschränkte divergente Folgen (siehe z.B. Abbildung 4.2.)
2. Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge. Falls $\lim_{n \rightarrow \infty} a_n = \infty$ oder $\lim_{n \rightarrow \infty} a_n = -\infty$. Dann ist für ein hinreichend grosses n_0 die Folge $\left(\frac{1}{a_n}\right)_{n \geq n_0}$ definiert, und es gilt: $\lim_{n \rightarrow \infty} \frac{1}{a_n} = 0$.

Beispiel 4.1.16 (Folgen a^n).

Für $0 < a \in \mathbb{R}$ gilt

$$\begin{aligned} \lim_{n \rightarrow \infty} a^n &= 0 && \text{für } a < 1, \\ \lim_{n \rightarrow \infty} a^n &= \infty && \text{für } a > 1. \end{aligned}$$

Beweis: Wir beweisen zunächst die zweite Aussage. Sei also $a > 1$, also $a = 1 + b$ mit $b > 0$. Wir können dann a^n mit Hilfe der Bernoulli-Ungleichung (Satz 1.4.5) nach unten abschätzen:

$$\begin{aligned} a^n &= (1 + b)^n \\ &\geq 1 + bn. \end{aligned}$$

Da die durch $b_n := 1 + bn$ definierte Folge nach oben unbeschränkt und eine **Minorante** der durch $a_n := a^n$ definierten Folge ist, geht $(a_n)_{n \in \mathbb{N}}$ gegen ∞ . Damit ist die zweite Aussage bewiesen.

Wenn $0 < a < 1$ dann ist $1 < \frac{1}{a}$. Nach der bereits bewiesenen zweiten Aussage gilt $\lim_{n \rightarrow \infty} \left(\frac{1}{a}\right)^n = \infty$, und aus Bemerkung 4.1.15.2 folgt dann Aussage 1. $\square$

4.2 Teilfolgen

Viele Folgen, denen wir begegnen, haben keinen Grenzwert. Manche oszillieren vielleicht, andere sind „chaotisch“, andere pendeln vielleicht zwischen verschiedenen Häufungspunkten (s. Definition 4.2.3). Was können wir trotzdem noch über solche Folgen sagen?

Beispiel 4.2.1 (Insulinspiegel).

Einem Versuchstier werde jede Stunde Blut entnommen und der Insulinspiegel (Insulinkonzentration) gemessen. Nach einigen Tagen ergibt sich das Bild in Abbildung 4.2. Man sieht, dass

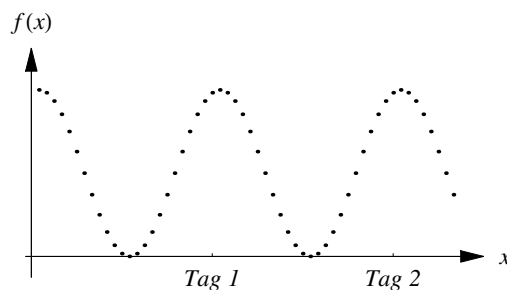


Abbildung 4.2: Die Insulinkonzentration schwankt periodisch.

immer wieder nach 24 Folgengliedern ein ähnlicher Wert angenommen wird.

Definition 4.2.2 (Teilfolge).

Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge und $n_0 < n_1 < n_2 < \dots$ eine aufsteigende Folge natürlicher Zahlen. Dann heißt die Folge

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

Teilfolge der Folge $(a_n)_{n \in \mathbb{N}}$.

Definition 4.2.3 (Häufungspunkt einer Folge).

Eine Zahl h heißt **Häufungspunkt** der Folge $(a_n)_{n \in \mathbb{N}}$, wenn es eine Teilfolge $(n_k)_{k \in \mathbb{N}}$ gibt, so dass die Folge $(a_{n_k})_{k \in \mathbb{N}}$ gegen h konvergiert.

Der folgende Satz, den wir hier nicht beweisen, liefert eine Charakterisierung von Häufungspunkten durch folgende zur Definition äquivalente Aussage: Es gibt Folgeglieder mit beliebig hohem Index, die beliebig nahe am Häufungspunkt liegen (Abstand kleiner als ein beliebig gewähltes positives ϵ).

Satz 4.2.4. Der Punkt h ist genau dann ein Häufungspunkt von $(a_n)_{n \in \mathbb{N}}$, wenn

$$\forall n \in \mathbb{N} \quad \forall \epsilon > 0 \quad \exists m \geq n \quad |a_m - h| < \epsilon.$$

4.2.1 *Der Satz von Bolzano-Weierstraß

Erstaunlich ist der folgende in der Mathematik sehr berühmte Satz:

Satz 4.2.5 (Bolzano-Weierstraß).

Jede beschränkte Folge $(a_n)_{n \in \mathbb{N}}$ reeller Zahlen besitzt eine konvergente Teilfolge (also einen Häufungspunkt).

Beweis: Da die Folge $(a_n)_{n \in \mathbb{N}}$ beschränkt ist, gibt es Zahlen $A, B \in \mathbb{R}$ mit

$$A \leq a_n \leq B \quad \forall n \in \mathbb{N}.$$

1. Schritt: Wir betrachten das Intervall

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

und konstruieren rekursiv eine Folge von Intervallen $[A_k, B_k]$, $k \in \mathbb{N}$, mit folgenden Eigenschaften:

1. In $[A_k, B_k]$ liegen unendlich viele Glieder der Folge (a_n) ,
2. $[A_k, B_k] \subset [A_{k-1}, B_{k-1}]$,
3. $B_k - A_k = 2^{-k}(B - A)$.

$k = 0$: Wir setzen $[A_0, B_0] := [A, B]$.

Wahl des Intervalls $[A_{k+1}, B_{k+1}]$ für $k > 0$: Sei das Intervall $[A_k, B_k]$ mit den Eigenschaften (1)-(3) bereits konstruiert. Sei $M := \frac{A_k + B_k}{2}$ die Mitte des Intervalls. Da in $[A_k, B_k]$ unendlich viele Glieder der Folge liegen, müssen in mindestens einem der Intervalle $[A_k, M]$ und $[M, B_k]$ unendlich viele Glieder der Folge liegen. Wir setzen

$$[A_{k+1}, B_{k+1}] := \begin{cases} [A_k, M], & \text{falls } [A_k, M] \text{ unendlich viele Folgenglieder hat,} \\ [M, B_k] & \text{sonst.} \end{cases}$$

Offenbar hat $[A_{k+1}, B_{k+1}]$ auch die Eigenschaften (1)-(3).

2. Schritt: Wir wählen eine Folge $(n_k)_{k \in \mathbb{N}}$ mit $a_{n_k} \in [A_k, B_k]$ für alle $k \in \mathbb{N}$. Für $k = 0$ setzen wir $n_0 = 0$. Sei nun $k \geq 1$. Da in dem Intervall $[A_k, B_k]$ unendlich viele Glieder der Folge $(a_n)_{n \in \mathbb{N}}$ liegen, können wir man ein $n_k > n_{k-1}$ mit $a_{n_k} \in [A_k, B_k]$ auswählen.

3. Schritt: Wir zeigen, dass die Teilfolge $(a_{n_k})_{k \in \mathbb{N}}$ konvergiert. Dann ist der Satz bewiesen. Es genügt zu zeigen, dass sie eine Cauchy-Folge ist (vgl. Definition 4.1.11 und Satz 4.1.12).

Sei $\epsilon > 0$ gegeben und ein $N \in \mathbb{N}$ so gewählt, dass die Länge des Intervalls $[A_N, B_N]$ durch $|B_N - A_N| = 2^{-N}(B - A) < \epsilon$ abgeschätzt wird. Dann gilt für alle $k, j \geq N$:

$$\begin{aligned} a_{n_k} &\in [A_k, B_k] \subset [A_N, B_N] \\ \text{und } a_{n_j} &\in [A_j, B_j] \subset [A_N, B_N]. \end{aligned}$$

Also ist

$$\begin{aligned} |a_{n_k} - a_{n_j}| &\leq |B_N - A_N| \\ &= 2^{-N}(B - A) < \epsilon. \end{aligned}$$

□

Beispiel 4.2.6 (Häufungspunkte von Folgen).

1. Die Folge $a_n = (-1)^n$ besitzt die Häufungspunkte $+1$ und -1 . Denn

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

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

$$\begin{aligned} \lim_{k \rightarrow \infty} a_{2k} &= \lim_{k \rightarrow \infty} \left(1 + \frac{1}{2k}\right) = 1 \quad \text{und analog} \\ \lim_{k \rightarrow \infty} a_{2k+1} &= -1. \end{aligned}$$

3. Die Folge $a_n = n$ besitzt keinen Häufungspunkt, da jede Teilfolge unbeschränkt ist.

4. Die Folge

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

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

5. Für jede konvergente Folge ist der Grenzwert ihr einziger Häufungspunkt.

4.2.2 *Limes inferior und Limes superior

Definition 4.2.7 (obere Schranke, untere Schranke, Supremum, Infimum).

Sei $A \subset \mathbb{R}$. Ein Element $s \in \mathbb{R}$ heißt **obere (untere) Schranke** von A , falls $a \leq s$ (bzw. $s \leq a$) $\forall a \in A$. Besitzt die Menge der oberen (unteren) Schranken von A ein Minimum s_1 (bzw. Maximum s_2), so heißt s_1 **Supremum** (bzw. heißt s_2 **Infimum**) von A .

Schreibweise:

$$\begin{aligned} \sup A &= s_1 \\ \inf A &= s_2. \end{aligned}$$

Also

$$\begin{aligned} \sup A &= \min\{s \in \mathbb{R} \mid s \text{ ist eine obere Schranke von } A\}, \\ \inf A &= \max\{s \in \mathbb{R} \mid s \text{ ist eine untere Schranke von } A\} \end{aligned}$$

Es sei nun $(x_n)_{n \in \mathbb{N}}$ eine beschränkte Folge in $\mathbb{R}$. Für jedes $n \in \mathbb{N}$ setzen wir

$$\begin{aligned} y_n &:= \sup(x_k)_{k \geq n} := \sup_{k \geq n} x_k := \sup\{x_k \mid k \geq n\}, \\ z_n &:= \inf(x_k)_{k \geq n} := \inf_{k \geq n} x_k := \inf\{x_k \mid k \geq n\}. \end{aligned}$$

Damit erhalten wir zwei neue Folgen. Offensichtlich ist $(y_n)_{n \in \mathbb{N}}$ eine *monoton fallende* und $(z_n)_{n \in \mathbb{N}}$ eine *monoton wachsende* Folge in $\mathbb{R}$. Deshalb existieren die Grenzwerte

$$\limsup_{n \rightarrow \infty} x_n := \overline{\lim}_{n \rightarrow \infty} x_n := \lim_{n \rightarrow \infty} (\sup_{k \geq n} x_k),$$

der **Limes superior**, und

$$\liminf_{n \rightarrow \infty} x_n := \underline{\lim}_{n \rightarrow \infty} x_n := \lim_{n \rightarrow \infty} (\inf_{k \geq n} x_k),$$

der **Limes inferior**.

Satz 4.2.8. Für eine konvergente Folge $(a_n)_{n \in \mathbb{N}}$ gilt

$$\lim_{n \rightarrow \infty} a_n = \limsup_{n \rightarrow \infty} a_n = \liminf_{n \rightarrow \infty} a_n. \quad (4.7)$$

□

4.3 Reihen

Kennen Sie Zenos Paradoxon vom Wettlauf des schnellsten Läufers der Antike, Achilles, mit einer Schildkröte, der vor dem Start ein kleiner Vorsprung gegeben wird? Die paradoxe Argumentation Zenos lautet: In dem Moment, wo Achilles an dem Ort s_0 ankommt, wo die Schildkröte gestartet ist, ist die Schildkröte selbst ja schon ein kleines Stückchen weitergekommen, sagen wir an die Stelle $s_1 > s_0$; Achilles muss also weiterlaufen, aber in dem Moment, wo er bei s_1 ankommt ist die Schildkröte wieder ein kleines Stückchen weitergekommen, sagen wir zum Punkt $s_2 > s_1$, usw. Der paradoxe Schluss Zenos ist, dass Achilles die Schildkröte nie einholen wird! Wie können wir dieses Paradoxon auflösen? Wir werden dies in Beispiel 4.3.16 erläutern, mit Hilfe des Begriffs der unendlichen Reihe, der das Thema dieses Abschnitts ist.

Definition 4.3.1 (Reihe).

Es sei $(a_k)_{k \in \mathbb{N}}$ eine Folge reeller Zahlen. Wir definieren eine neue Folge s_n durch

$$s_n := \sum_{k=0}^n a_k, \quad n \in \mathbb{N}.$$

Die Folge $(s_n)_{n \in \mathbb{N}}$ heißt **Reihe**, sie wird mit $\sum_k a_k$ bezeichnet und s_n heißt die **n -te Partialsumme**.

Die ersten vier Partialsummen sind:

$$\begin{aligned} s_0 &= a_0, \\ s_1 &= a_0 + a_1, \\ s_2 &= a_0 + a_1 + a_2, \\ s_3 &= a_0 + a_1 + a_2 + a_3, \\ s_4 &= a_0 + a_1 + a_2 + a_3 + a_4. \end{aligned}$$

Bemerkung 4.3.2 (Beziehung zwischen Folgen und Reihen).

Wir haben zu jeder Folge eine Reihe definiert, und zwar durch

$$s_0 := a_0, \quad s_{n+1} = s_n + a_n, \quad n \in \mathbb{N}.$$

Diese Beziehung lässt sich offensichtlich auch umkehren, d.h. zu jeder Reihe $(s_n)_{n \in \mathbb{N}}$ gibt es eine entsprechende Folge $(a_k)_{k \in \mathbb{N}}$ von Summanden:

$$a_0 := s_0, \quad a_n = s_{n+1} - s_n, \quad n \in \mathbb{N}.$$

Beispiel 4.3.3 (für Reihen).

1. Die **harmonische Reihe** $\sum_{k=1}^{\infty} \frac{1}{k}$ divergiert.

Denn $|s_{2n} - s_n| = \sum_{k=n+1}^{2n} \frac{1}{k} \geq \frac{n}{2n} = \frac{1}{2}$, also ist $(s_n)_{n \in \mathbb{N}}$ keine Cauchy-Folge und divergiert deshalb. Es gilt

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

2. Die Reihe $\sum_{k=1}^{\infty} \frac{1}{k^2}$ konvergiert. Offensichtlich ist die Folge der Partialsummen $(s_n)_{n \geq 1}$ monoton wachsend. Desweiteren gilt

$$\begin{aligned} s_n &= \sum_{k=1}^n \frac{1}{k^2} \\ &\leq 1 + \sum_{k=2}^n \frac{1}{k(k-1)} \\ &= 1 + \sum_{k=2}^n \left(\frac{1}{(k-1)} - \frac{1}{k} \right) \\ &= 1 + 1 - \frac{1}{n} < 2, \end{aligned}$$

also ist $(s_n)_{n \in \mathbb{N}}$ beschränkt und konvergiert daher nach Satz 4.1.12.2.

3. Die **geometrische Reihe** $\sum_{k=0}^{\infty} c^k$ mit $0 < |c| < 1$ konvergiert gegen $\frac{1}{1-c}$, denn $\sum_{k=0}^n c^k = \frac{1-c^{n+1}}{1-c}$, wie man leicht zeigen kann, und $\lim_{n \rightarrow \infty} c^{n+1} = 0$.

Satz 4.3.4 (Rechenregeln für konvergente Reihen).

Es seien $\sum_k a_k$ und $\sum_k b_k$ konvergente Reihen, sowie $\alpha \in \mathbb{R}$. Dann gilt:

1. Die Reihe $\sum_k (a_k + b_k)$ konvergiert und

$$\sum_{k=0}^{\infty} (a_k + b_k) = \sum_{k=0}^{\infty} a_k + \sum_{k=0}^{\infty} b_k.$$

2. Die Reihe $\sum_k (\alpha a_k)$ konvergiert und

$$\sum_{k=0}^{\infty} (\alpha a_k) = \alpha \sum_{k=0}^{\infty} a_k.$$

4.3.1 Konvergenzkriterien für Reihen**Satz 4.3.5** (Cauchy-Kriterium).

Die folgenden zwei Aussagen sind einander äquivalent:

1. $\sum_k a_k$ ist konvergent.
2. $\forall \epsilon > 0 \quad \exists N \in \mathbb{N} \quad \forall m, n \text{ mit } N \leq n < m :$

$$\left| \sum_{k=n+1}^m a_k \right| < \epsilon$$

Beweis: Es gilt $s_m - s_n = \sum_{k=n+1}^m a_k$ für $m > n$. Somit ist $(s_n)_{n \in \mathbb{N}}$ genau dann eine Cauchy-Folge und somit genau dann konvergent, wenn (2.) wahr ist. $\square$

Satz 4.3.6 (Kovergenz monotoner beschränkter Reihen).

Es sei $\sum_k a_k$ eine Reihe mit $a_k > 0$, $k \in \mathbb{N}$. Dann ist $\sum_k a_k$ genau dann konvergent, wenn $(s_n)_{n \in \mathbb{N}}$ beschränkt ist. Die Reihe konvergiert gegen $\sup_{n \in \mathbb{N}} s_n$.

Beweis: Die Folge $(s_n)_{n \in \mathbb{N}}$ der Partialsummen ist monoton wachsend und konvergiert nach Satz 4.1.12.2, wenn sie (s_n) beschränkt ist. Da die Beschränktheit eine notwendige Bedingung für Konvergenz ist, folgt aus Satz 4.1.9.1. Die kleinste Zahl welche größer oder gleich allen s_n ist, ist $\sup_{n \in \mathbb{N}} s_n$. Die Konvergenz der Reihe gegen diese Zahl folgt aus Satz 4.2.8, wobei wir dies hier nicht im Detail begründen. $\square$

4.3.2 *Alternierende Reihen

In diesem Teilabschnitt betrachten wir nur Reihen $\sum_k a_k$ mit nicht-negativen Summanden, d.h. $a_k \geq 0 \quad \forall k \in \mathbb{N}$.

Satz 4.3.7 (Leibnizsches Kriterium).

Es sei $(a_k)_{k \in \mathbb{N}}$ eine fallende Nullfolge. Dann konvergiert $\sum_k (-1)^k a_k$.

Beweis: Die Folge $(s_{2n})_{n \in \mathbb{N}}$ (gerade Indizes) ist wegen

$$s_{2n+2} - s_{2n} = -a_{2n+1} + a_{2n+2} \leq 0, \quad n \in \mathbb{N}$$

monoton fallend. Analog ist $(s_{2n+1})_{n \in \mathbb{N}}$ wegen

$$s_{2n+3} - s_{2n+1} = a_{2n+2} - a_{2n+3} \geq 0, \quad n \in \mathbb{N}$$

monoton wachsend. Desweiteren ist $s_{2n+1} \leq s_{2n}$, und somit

$$s_{2n+1} \leq s_0 \text{ und } s_{2n} \geq s_1, \quad n \in \mathbb{N}$$

Wegen ihrer Beschränktheit konvergieren diese Teilfolgen, also

$$\lim_{n \rightarrow \infty} s_{2n} = \gamma, \quad \lim_{n \rightarrow \infty} s_{2n+1} = \delta$$

Daher ist

$$\gamma - \delta = \lim_{n \rightarrow \infty} (s_{2n} - s_{2n+1}) = \lim_{n \rightarrow \infty} a_{2n+1} = 0.$$

Daher gibt es $\epsilon > 0$, $N_1, N_2 \in \mathbb{N}$ mit

$$\begin{aligned} |s_{2n} - \gamma| &< \epsilon, & \text{für } 2n \geq N_1 \text{ und} \\ |s_{2n+1} - \gamma| &< \epsilon, & \text{für } 2n+1 \geq N_2. \end{aligned}$$

Somit gilt $|s_n - \gamma| < \epsilon$ für $n \geq \max(N_1, N_2)$ und die Konvergenz von $(s_n)_{n \in \mathbb{N}}$ ist gezeigt. $\square$

Beispiel 4.3.8 (alternierende harmonische Reihe).

Die **alternierende harmonische Reihe**

$$\sum_{k=1}^{\infty} \frac{(-1)^{k+1}}{k} = 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \frac{1}{5} - \dots$$

konvergiert.

4.3.3 *Absolute Konvergenz

Definition 4.3.9 (absolute Konvergenz).

Eine Reihe $\sum_k a_k$ heißt **absolut konvergent**, falls $\sum_k |a_k|$ konvergiert.

Satz 4.3.10 (Aus absoluter Konvergenz folgt Konvergenz.).

Jede absolut konvergente Reihe konvergiert.

Beweis: Sei $\sum a_k$ absolut konvergent, d.h. $\sum |a_k|$ konvergiere. Dann gilt das Cauchy-Kriterium:

$$\forall \epsilon > 0 \quad \exists N : \quad \sum_{k=n+1}^m |a_k| < \epsilon \quad \text{für } m > n \geq N.$$

Wegen

$$\left| \sum_{k=n+1}^m a_k \right| \leq \sum_{k=n+1}^m |a_k| < \epsilon \quad \text{für } m > n \geq N$$

folgt, dass $\sum a_k$ konvergiert. $\square$

Definition 4.3.11 (bedingte Konvergenz).

Die Reihe $\sum a_k$ heißt **bedingt konvergent**, falls $\sum_k a_k$ konvergiert, aber $\sum_k |a_k|$ nicht konvergiert.

Lemma 4.3.12. (Dreiecksungleichung für absolut konvergente Reihen)

Für jede absolut konvergente Reihe $\sum a_k$ gilt die **verallgemeinerte Dreiecksungleichung**

$$\left| \sum_{k=0}^{\infty} a_k \right| \leq \sum_{k=0}^{\infty} |a_k|. \quad (4.8)$$

Beweis: Sei $\epsilon > 0$ beliebig und N so gewählt, dass

$$\sum_{k=N+1}^{\infty} |a_k| < \epsilon. \quad (4.9)$$

Dann gilt

$$\left| \sum_{k=0}^{\infty} a_k \right| = \left| \sum_{k=0}^N a_k + \sum_{k=N+1}^{\infty} a_k \right| \quad (4.10)$$

$$\leq \left| \sum_{k=0}^N a_k \right| + \left| \sum_{k=N+1}^{\infty} a_k \right| \quad (4.11)$$

$$\leq \sum_{k=0}^N |a_k| + \epsilon \quad (4.12)$$

$$\leq \sum_{k=0}^{\infty} |a_k| + \epsilon.$$

Dabei haben wir im Schritt von (4.10) nach (4.11) die Dreiecksungleichung für reelle Zahlen, im Schritt von (4.11) nach (4.12) zur Abschätzung des ersten Summanden die Dreiecksungleichung für Summen endlich vieler reeller Zahlen sowie die Abschätzung (4.9) verwendet. Insgesamt erhalten wir also

$$\left| \sum_{k=0}^{\infty} a_k \right| \leq \sum_{k=0}^{\infty} |a_k| + \epsilon.$$

für beliebig kleine $\epsilon > 0$. Daraus folgt (4.8). $\square$

Definition 4.3.13 (Majorante und Minorante einer Reihe).

Seien $\sum a_k$ und $\sum b_k$ Reihen und es gelte $b_k \geq 0 \quad \forall k \in \mathbb{N}$. Dann heißt die Reihe $\sum b_k$ **Majorante** bzw. **Minorante** von $\sum a_k$, falls es ein $k_0 \in \mathbb{N}$ gibt mit

$$|a_k| \leq b_k \quad \text{bzw.} \quad |a_k| \geq b_k \quad \text{für alle } k \geq k_0.$$

Satz 4.3.14 (Majorantenkriterium).

Besitzt eine Reihe eine konvergente Majorante, so konvergiert sie absolut.

Beweis: Es sei $\sum a_k$ eine Reihe und $\sum b_k$ eine konvergente Majorante. Dann gibt es ein k_0 mit $|a_k| \leq b_k$ für $k \geq k_0$. Nach Satz (4.3.5) gibt es zu $\epsilon > 0$ ein $N \geq k_0$ mit $\sum_{k=n+1}^m b_k < \epsilon$ für $m > n \geq N$. Da $\sum b_k$ eine Majorante für $\sum a_k$ ist, erhalten wir

$$\sum_{k=n+1}^m |a_k| \leq \sum_{k=n+1}^m b_k < \epsilon \quad \text{für } m > n \geq N.$$

Nach Satz (4.3.5) konvergiert $\sum |a_k|$, dass heißt $\sum a_k$ konvergiert absolut. $\square$

Beispiel 4.3.15 (Majorisierung der geometrischen Reihe).

$\sum_{k=1}^{\infty} \frac{1}{k^m}$, $m \geq 2$ konvergiert. Eine konvergente Majorante ist $\sum_{k=1}^{\infty} \frac{1}{k^2}$, siehe Beispiel 4.3.3.2.

Beispiel 4.3.16 (Achilles und die Schildkröte). }

Wir werden nun Zenos Paradoxon vom Wettlauf zwischen Achilles und der Schildkröte auflösen. Sagen wir, Achilles ist c -mal schneller als die Schildkröte, und die Schildkröte startet am Ort s_0 , mit $c > 1$ und $s_0 > 0$. Wir wollen mit Hilfe einer Reihe den Ort berechnen, an dem Achilles die Schildkröte einholt. Dafür betrachten wir die Wegstücke zwischen den Stellen s_i aus Zenos Argumentation, an denen die Schildkröte immer wieder ein Stück weiter ist als Achilles, wenn er gerade bei s_{i-1} ankommt. Während Achilles das neue Stück $s_i - s_{i-1}$ läuft, schafft die Schildkröte nur ein c -tel der Entfernung, also $s_{i+1} - s_i = (s_i - s_{i-1})/c$. Daraus (und aus der Tatsache, dass $s_1 - s_0 = s_0/c$) können wir induktiv schliessen, dass

$$s_i - s_{i-1} = s_0 \frac{1}{c^i} \quad \text{also} \quad s_k = s_0 \sum_{i=0}^k \frac{1}{c^i},$$

und wir erkennen, dass wir es hier mit einer geometrischen Reihe zu tun haben, deren Grenzwert wir kennen! Achilles überholt die Schildkröte genau am Ort

$$s_0 \sum_{i=0}^{\infty} \frac{1}{c^i} = s_0 \frac{1}{1 - \frac{1}{c}} = \frac{s_0 c}{c - 1}.$$

4.4 Exponentialfunktion und Logarithmus

Für jedes $x \in \mathbb{R}$ definieren wir die Exponentialfunktion durch die folgende Reihe:

$$\exp(x) := \sum_{k=0}^{\infty} \frac{x^k}{k!}, \quad (4.13)$$

Diese Funktion wird Ihnen in Ihrem Studium und in der Praxis noch häufig begegnen – sie spielt eine äußerst wichtige Rolle in vielen praktischen Anwendungen, und es lohnt sich, sich mit ihren Eigenschaften gut vertraut zu machen.

4.4.1 Eigenschaften der Exponentialfunktion

Gehen wir zunächst in die Finanzmathematik. Bei jährlicher Verzinsung mit Zinssatz p wächst ein Anfangskapital K nach m Jahren auf

$$K_m = K \left(1 + \frac{p}{100}\right)^m.$$

Bei unterjährlicher Verzinsung, wobei das Jahr in n Zinsperioden unterteilt ist, wächst das Startkapital nach einem Jahr auf

$$K_1^{(n)} = K \left(1 + \frac{p}{100n}\right)^n.$$

Nach m Jahren ergibt sich bei der gleichen unterjährigen Verzinsung ein Kapital von

$$K_m^{(n)} = K \left(1 + \frac{p}{100n}\right)^{mn}.$$

Wählen wir feste Parameter $m = 1$, $K = 1$ und $x = \frac{p}{100}$, und lassen die Zinsperioden immer kleiner werden ($n \rightarrow \infty$), so ergibt sich als Grenzwert

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

Insbesondere gilt somit

$$\exp(1) = e,$$

wobei e die Eulersche Zahl aus Beispiel 4.1.13 ist. Wir schreiben auch e^x anstatt $\exp(x)$. *Ausblick:* Die Exponentialfunktion erfüllt auch (ist Lösung von) der **gewöhnlichen Differentialgleichung** (genauer: des **Anfangswertproblems** mit Anfangswert x_0)

$$\begin{cases} \frac{d}{dt} x(t) &= a \cdot x(t), \\ x(0) &= x_0. \end{cases} \quad (4.14)$$

Die Lösung des Anfangswertproblems ist $x(t) = x_0 e^{at} = x_0 \exp(at)$.

Satz 4.4.1 (Eigenschaften der Exponentialfunktion).

1. $\exp(x + y) = \exp(x) \cdot \exp(y) \quad \forall x, y \in \mathbb{R}.$
2. $1 + x \leq \exp(x) \quad \forall x \in \mathbb{R}.$
3. $\exp(x) \leq \frac{1}{1-x} \quad \forall x < 1.$
4. $\exp(x)$ ist streng monoton wachsend.
5. Das Bild von $\exp(x)$ ist $\mathbb{R}^+$.

Wir werden weiter unten nur Eigenschaft (1.) beweisen, und zwar unter Benutzung des folgenden Satzes.

***Satz 4.4.2** (Cauchy-Produkt von absolut konvergenten Reihen).

Falls $\sum_j a_j$ und $\sum_k b_k$ absolut konvergieren, so konvergiert auch $\sum_n \sum_{k=0}^n a_k b_{n-k}$ absolut und

$$\left(\sum_{j=0}^{\infty} a_j\right) \left(\sum_{k=0}^{\infty} b_k\right) = \sum_{n=0}^{\infty} \sum_{k=0}^n a_k b_{n-k} \quad (\text{Cauchy-Produkt}) \quad (4.15)$$

Zu zeigen ist also, daß $\sum_{k=0}^{\infty} \frac{x^k}{k!}$ ist für jedes $x \in \mathbb{R}$ absolut konvergent ist. Dazu benutzen wir das **Quotientenkriterium**.

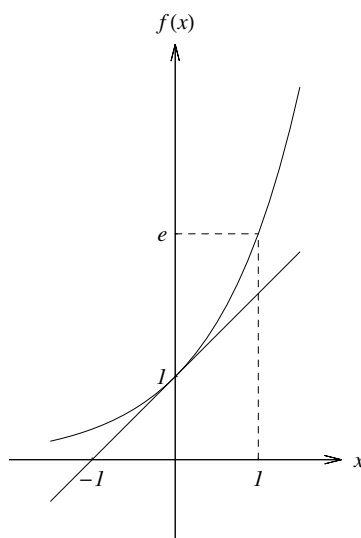


Abbildung 4.3: Die Exponentialfunktion

Satz 4.4.3 (Quotientenkriterium für absolute Konvergenz von Reihen).

Sei $\sum_k a_k$ eine Reihe mit $a_n \neq 0 \quad \forall n \geq N$. Es gebe eine reelle Zahl θ mit $0 < \theta < 1$, so dass

$$\left| \frac{a_{k+1}}{a_k} \right| \leq \theta \quad \forall k \geq N.$$

Dann konvergiert $\sum_k a_k$ absolut.

Beweis von Theorem 4.4.1: Wir weisen nur Eigenschaft (1.) nach. Für die Exponentialreihe gilt für $k \geq 2|x|$:

$$\left| \frac{\frac{x^{k+1}}{(k+1)!}}{\frac{x^k}{k!}} \right| = \frac{|x|}{k+1} \leq \frac{1}{2},$$

d.h. sie konvergiert absolut für jedes $x \in \mathbb{R}$. Daher existiert ihr Cauchy-Produkt und wir erhalten

$$\begin{aligned} \exp(x) \cdot \exp(y) &= \left(\sum_{j=0}^{\infty} \frac{x^j}{j!} \right) \left(\sum_{k=0}^{\infty} \frac{y^k}{k!} \right) \\ &= \sum_{n=0}^{\infty} \left(\sum_{k=0}^n \frac{x^k}{k!} \frac{y^{n-k}}{(n-k)!} \right). \end{aligned}$$

Unter Verwendung des binomischen Lehrsatzes 1.5 1.5.1 machen wir folgende Nebenrechnung.

$$\begin{aligned} \sum_{k=0}^n \frac{x^k}{k!} \frac{y^{n-k}}{(n-k)!} &= \frac{1}{n!} \sum_{k=0}^n \frac{n!}{k!(n-k)!} x^k y^{n-k} \\ &= \frac{1}{n!} \sum_{k=0}^n \binom{n}{k} x^k y^{n-k} \\ &= \frac{1}{n!} (x+y)^n. \end{aligned}$$

Somit erhalten wir

$$\exp(x) \cdot \exp(y) = \sum_{n=0}^{\infty} \frac{(x+y)^n}{n!} = \exp(x+y).$$

□

4.4.2 Der natürliche Logarithmus

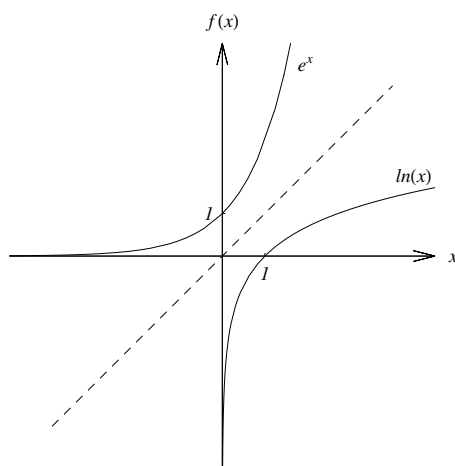


Abbildung 4.4: Die natürliche Logarithmusfunktion und die Exponentialfunktion sind zueinander invers.

Die Exponentialfunktion steigt streng monoton und jeder Wert $y > 0$ wird genau einmal von e^x angenommen. Deshalb können wir die Umkehrfunktion definieren, die wir den **natürlichen Logarithmus** nennen, und mit dem Symbol $\ln(x)$ bezeichnen:

$$\begin{aligned} \ln : \mathbb{R}_+ &\longrightarrow \mathbb{R}, \\ x &\longmapsto \ln(x). \end{aligned}$$

Es gilt nach Definition

$$\ln(e^x) = x \quad \forall x \in \mathbb{R}.$$

In Abbildung 4.4 veranschaulichen wir, wie der Graph der natürlichen Logarithmusfunktion durch Spiegelung an der Diagonalen aus dem Graph der Exponentialfunktion hervorgeht. Man beachte, dass der Logarithmus nur für positive Argumente definiert ist, weil die Exponentialfunktion nur positive Werte annehmen kann. Eine genauere Betrachtung des Logarithmus als Umkehrfunktion zur Exponentialfunktion erfolgt in Beispiel 1 in Kapitel 4.6.

4.4.3 Potenzen und Logarithmen zu einer positiven Basis

Statt e^x können wir auch b^x , $b > 0$ bilden. Wir definieren

$$b^x := \exp(x \ln(b)). \quad (4.16)$$

Die Funktion $x \mapsto b^x$, $x \in \mathbb{R}$, heißt **Exponentialfunktion zur Basis b** . Für $b \neq 1$ existiert auch die Umkehrfunktion zu b^x . Sie wird **Logarithmus zur Basis b** genannt und mit

$$x \mapsto \log_b(x), \quad x \in \mathbb{R}^+ \quad (4.17)$$

bezeichnet. Es gilt

$$\log_b(x) = \frac{\ln(x)}{\ln(b)}, \quad (4.18)$$

denn aus $x = b^y = \exp(y \log(b))$ folgt $\ln(x) = y \ln(b) = \log_b(x) \log(b)$.

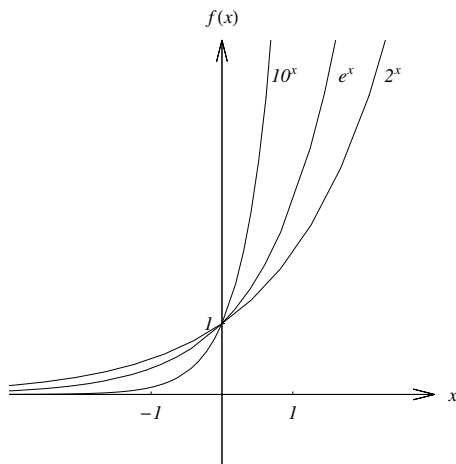


Abbildung 4.5: Die wichtigsten Exponentialfunktionen, zur Basis 2, e und 10.

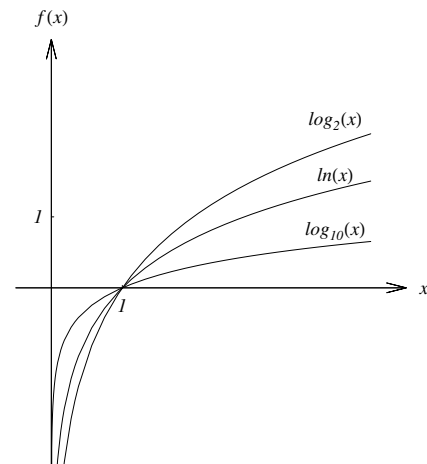


Abbildung 4.6: Die wichtigsten Logarithmusfunktionen, zur Basis 2, e und 10.

4.5 Stetigkeit

Im Folgenden bezeichnet U immer eine nichtleere Teilmenge von $\mathbb{R}$, also z.B. $U = \mathbb{R}$, $U = (a, b)$, $U = [a, b]$, $U = [\infty, 0]$ etc. Wir betrachten reellwertige Funktionen mit Definitionsmenge

U :

$$\begin{aligned} f : U &\rightarrow \mathbb{R} \\ x &\mapsto f(x), \end{aligned}$$

also $x \in U$, $f(x) \in \mathbb{R}$. Die **Wertemenge** von f ist definiert als

$$f(U) := \{y \in \mathbb{R} : \exists x \in U \quad f(x) = y\}.$$

Wir wollen uns jetzt mit der allgemeinen (unpräzisen) Frage beschäftigen: Wie ändert sich der Funktionswert, wenn das Argument „ein bißchen“ geändert wird?

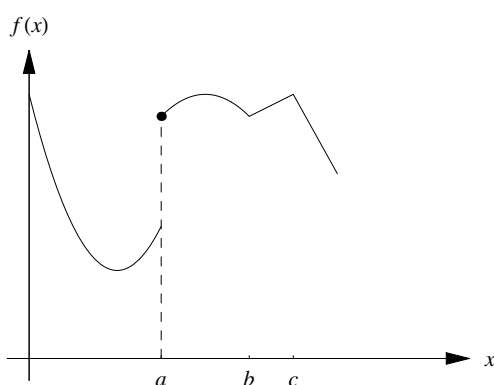


Abbildung 4.7: Graph einer unstetigen Funktion

Beispiel 4.5.1 (einer nicht-stetigen Funktion).

Wir betrachten die Funktion

$$\begin{aligned} f : \mathbb{R} &\rightarrow \mathbb{R} \\ x &\mapsto f(x) := \begin{cases} -1 & \text{für } x < 0, \\ 1 & \text{für } x \geq 0. \end{cases} \end{aligned}$$

Die Funktion f „macht einen Sprung“ bei $x = 0$. Genauer: Es gilt $f(0) = 1$, aber $f(-\epsilon) = -1$, für alle $\epsilon > 0$. Je nachdem, von welcher Seite sich eine monotone Folge $(x^{(n)})_{n \in \mathbb{N}}$ dem Grenzwert $x = 0$ nähert, entweder von links oder von rechts, hat die Folge der Bilder $f(x^{(n)})$ unterschiedliche Grenzwerte. Wir werden eine Eigenschaft von Funktionen definieren, bei denen der Grenzwert jeweils eindeutig ist (also nicht von der speziellen Folge der Argumente abhängt). In Abbildung 4.7 zeigen wir ein weiteres Beispiel einer unstetigen Funktion.

Zunächst eine Notation:

Definition 4.5.2 (Grenzwert einer Funktion).

Seien $f : U \rightarrow \mathbb{R}$ und $x_0 \in \bar{U}$. Dabei bezeichnen wir mit $\bar{U}$ den **Abschluß** von U , d.h. die

Menge aller Punkte in $\mathbb{R}$, die durch eine Folge von Punkten in U approximiert werden können, also Grenzwert einer dieser Folge sind. Wir schreiben

$$\lim_{x \rightarrow x_0} f(x) = y, \quad (4.19)$$

falls für jede Folge $(x^{(n)})_{n \in \mathbb{N}}$ mit $x^{(n)} \in U$ und $\lim_{n \rightarrow \infty} x^{(n)} = x_0$ die Folge der Bilder $f(x^{(n)})$ gegen y konvergiert, d.h. $\lim_{n \rightarrow \infty} f(x^{(n)}) = y$.

Bemerkung 4.5.3. Falls $x_0 \in U$ und Eigenschaft (4.19) gilt, dann ist der Grenzwert $y = f(x_0)$, da durch $x^{(n)} = x_0$ offensichtlich eine Folge mit Grenzwert x_0 definiert ist.

Definition 4.5.4 (Folgenkriterium für die Stetigkeit einer Funktion).

1. Eine Funktion $f : U \rightarrow \mathbb{R}$ heißt **stetig in** $x_0 \in U$, wenn

$$\lim_{x \rightarrow x_0} f(x) = f(x_0).$$

2. Sei $V \subset U$. Eine Funktion $f : U \rightarrow \mathbb{R}$ heißt **stetig in** V (auf V), wenn f in jedem Punkt von V stetig ist.

Beispiel 4.5.5. 1. Die Funktion f aus (4.5.1) ist stetig in $\mathbb{R} \setminus \{0\}$, aber sie ist nicht stetig in $x = 0$.

2. Sei $c \in \mathbb{R}$ und sei $f : \mathbb{R} \rightarrow \mathbb{R}$ definiert durch $f(x) = c$ (konstante Funktion). Dann ist f stetig auf $\mathbb{R}$.

3. Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$, definiert durch $f(x) = x$ ist stetig auf $\mathbb{R}$.

Beweis dazu: Sei $\lim_{n \rightarrow \infty} x^{(n)} = x_0$. Dann gilt nach Definition von f :

$$\lim_{n \rightarrow \infty} f(x^{(n)}) = \lim_{n \rightarrow \infty} x^{(n)} = x_0.$$

□

Satz 4.5.6 (Addition, Multiplikation und Division stetiger Funktionen).

Seien $f, g : U \rightarrow \mathbb{R}$ auf U stetige Funktionen. Dann gilt:

1. $f + g$ ist stetig auf U .
2. $f \cdot g$ ist stetig auf U .
3. Sei zusätzlich $f(x) \neq 0$ für alle $x \in U$. Dann ist die durch $\frac{1}{f(x)}$ definierte Funktion stetig auf U .

Beweis: Der Beweis folgt aus dem entsprechenden Satz für Folgen (Satz 4.1.9).

□

Bemerkung 4.5.7. Aus (2) folgt insbesondere, dass mit f auch $-f$ stetig ist. (Nimm $g = -1$.) Wegen (1) folgt auch die Stetigkeit von $f - g$. Unter der Bedingung von (3) folgt die Stetigkeit von $\frac{g}{f}$.

Satz 4.5.8 (Komposition stetiger Funktionen).

Seien $g : U \rightarrow \mathbb{R}$ und $f : V \rightarrow \mathbb{R}$ stetig und $g(U) \subset V$. Dann ist die *Komposition (Verknüpfung)* $f \circ g : U \rightarrow \mathbb{R}$, definiert durch $(f \circ g)(x) = f(g(x))$, stetig.

Beweis: Zum Beweis der Stetigkeit in $x_0 \in U$, sei $\lim_{n \rightarrow \infty} x^{(n)} = x_0$. Dann gilt wegen der Stetigkeit von g , dass $\lim_{n \rightarrow \infty} g(x^{(n)}) = g(x_0)$, und somit wegen der Stetigkeit von f in $g(x_0)$ auch

$$\begin{aligned} \lim_{n \rightarrow \infty} (f \circ g)(x^{(n)}) &= \lim_{n \rightarrow \infty} f(g(x^{(n)})) \\ &= f\left(\lim_{n \rightarrow \infty} g(x^{(n)})\right) \\ &= f(g(x_0)) \\ &= (f \circ g)(x_0). \end{aligned}$$

□

Beispiel 4.5.9 (Wichtige stetige Funktionen).

1. Polynome sind stetige Funktionen: $p(x) = \sum_{k=0}^n a_k x^k$. Nach Beispiel (4.5.5.1) ist $x \mapsto x$ stetig, wegen Satz (4.5.6.2) ist $x \mapsto \underbrace{x \cdot \dots \cdot x}_{k \text{ mal}} = x^k$ stetig und wegen Satz (4.5.6.1) ist p stetig.
2. Die Exponentialfunktion $e^x = \sum_{n=0}^{\infty} \frac{x^n}{n!}$ ist stetig auf $\mathbb{R}$. Ebenso sind $\sin x$, $\cos x$ stetig.
3. Die Funktion $f : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R}$, definiert durch $f(x) = \frac{1}{x}$, ist stetig.
4. (Verallgemeinerung von (3)) Gebrochen-rationale Funktionen lassen sich darstellen als $f(x) = \frac{p(x)}{q(x)}$, wobei $p(x)$ und $q(x)$ Polynome sind, und q ist *nicht* das Nullpolynom ist. Dann hat q endlich viele reelle Nullstellen $x_1, \dots, x_N$ (und evtl. auch nicht reelle) und

$$f : \mathbb{R} \setminus \{x_1, \dots, x_N\} \rightarrow \mathbb{R}$$

ist stetig.

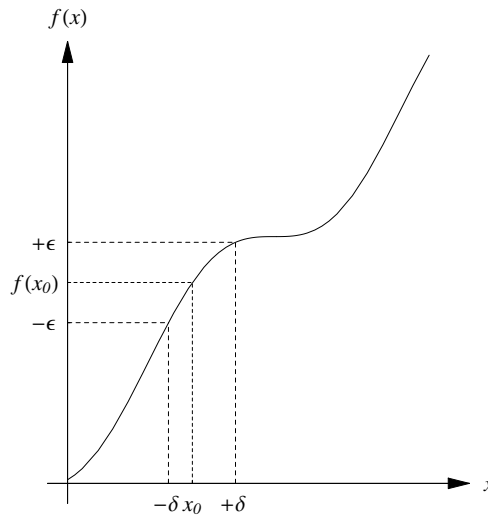
Eine nützliche äquivalente (alternative) Stetigkeitsdefinition ist durch die δ - ϵ -Definition gegeben.

Satz 4.5.10 (δ - ϵ -Kriterium für Stetigkeit).

Sei $f : U \rightarrow \mathbb{R}$. Äquivalent zur Stetigkeit von f in $x_0 \in U$ ist die Aussage:

$$\forall \epsilon > 0 \quad \exists \delta > 0 \quad \forall x \in U : \quad |x_0 - x| < \delta \Rightarrow |f(x_0) - f(x)| < \epsilon.$$

(Siehe auch Abbildung 4.8)

Abbildung 4.8: Illustration zum $\delta - \epsilon$ -Kriterium

Beispiel 4.5.11 (Stetigkeit von $f(x) = \frac{1}{x}$ in $x_0 \neq 0$).

1. Seien $f(x) = \frac{1}{x}$, $x_0 = 5$, $\epsilon = \frac{1}{10}$ vorgegeben. Es gilt

$$|f(x) - f(5)| = \left| \frac{1}{x} - \frac{1}{5} \right| = \left| \frac{5 - x}{5x} \right| \stackrel{!}{<} \frac{1}{10}.$$

Wähle $\delta = 1$, dann gilt $4 < x < 6$, $20 < 5x < 30$, $-1 < 5 - x < 1$, also

$$\left| \frac{5 - x}{5x} \right| \leq \frac{1}{20} < \frac{1}{10}.$$

Also ist die δ - ϵ -Bedingung für f für $x_0 = 5$, $\epsilon = \frac{1}{10}$ z.B. mit $\delta = 1$ erfüllt.

2. Allgemein sei nun $x_0 > 0$, und $\epsilon > 0$.

Unter der Bedingung $\delta < \frac{1}{2}x_0$ gilt $x \in (x_0 - \frac{1}{2}x_0, x_0 + \frac{1}{2}x_0) = (\frac{1}{2}x_0, \frac{3}{2}x_0)$. Und somit

$$|f(x) - f(x_0)| = \frac{|x_0 - x|}{|x \cdot x_0|} < \frac{\delta}{\frac{1}{2}x_0 \cdot x_0} = \frac{2\delta}{x_0^2} \stackrel{!}{<} \epsilon.$$

Wähle also $\delta < \min\{\frac{\epsilon x_0^2}{2}, \frac{1}{2}x_0\}$. Dann ist die geforderte Bedingung erfüllt.

Die Wahl ist im Fall $x_0 < 0$ analog: $\delta = \min\{\frac{\epsilon x_0^2}{2}, \frac{1}{2}|x_0|\}$.

Im Folgenden sollte klar werden, warum die Stetigkeit einer Funktion eine so nützliche Eigenschaft ist.

Satz 4.5.12 (Nullstellensatz und Zwischenwertsatz).

1. **(Nullstellensatz)** Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig und $f(a) < 0 < f(b)$ (bzw. $f(a) > 0 > f(b)$). Dann hat f in $]a, b[$ mindestens eine Nullstelle.
2. **(Zwischenwertsatz)** Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig mit $f(a) < f(b)$ (bzw. $f(a) > f(b)$). Dann nimmt f auf $[a, b]$ jeden Wert des Intervalls $[f(a), f(b)]$ (bzw. $[f(b), f(a)]$) an.

Beweis: Zu (2): Benutze (1).

Zu (1): Definiere eine Intervallschachtelung. Seien (ohne Einschränkung der Allgemeinheit) $f(a) < 0$, $f(b) > 0$. Wir definieren

$$\begin{aligned} [x_l^{(0)}, x_r^{(0)}] &:= [a, b], \\ x^{(i)} &:= \frac{x_l^{(i)} + x_r^{(i)}}{2} \quad \text{für alle } i \in \mathbb{N}. \end{aligned}$$

Falls $f(x^{(i)}) < 0$, so definieren wir

$$[x_l^{(i+1)}, x_r^{(i+1)}] = [x^{(i)}, x_r^{(i+1)}].$$

Falls $f(x^{(i)}) > 0$, so definieren wir

$$[x_l^{(i+1)}, x_r^{(i+1)}] = [x_l^{(i)}, x^{(i)}].$$

Und falls $f(x^{(i)}) = 0$, dann ist eine Nullstelle gefunden.

Falls keines der $x^{(i)}$ eine Nullstelle ist, so definiert die Intervallschachtelung $[x^{(0)}, x_r^{(0)}] \subset [x^{(1)}, x_r^{(1)}] \subset \dots$ eine reelle Zahl, die Nullstelle von f ist.

Denn sei x_0 diese Zahl. Wegen $\lim_{i \rightarrow \infty} x_l^{(i)} = x_0$ und der Stetigkeit von f gilt $\lim_{i \rightarrow \infty} f(x_l^{(i)}) = f(x_0)$, und wegen $f(x_l^{(i)}) < 0 \quad \forall i = N$, ist $f(x_0)$ Grenzwert einer Folge negativer Zahlen, kann also nicht positiv sein. Analog zeigt man, dass $f(x_0)$ nicht negativ ist. Es folgt $f(x_0) = 0$. $\square$

Bemerkung 4.5.13. Satz 4.5.12 garantiert die *Existenz* einer Nullstelle unter bestimmten Bedingungen. Die Intervallschachtelung (siehe Abbildung 4.9) gibt ein mögliches Verfahren zur Approximation einer Nullstelle an.

Definition 4.5.14 (globale und lokale Extrema einer Funktion).

Seien $f : U \rightarrow \mathbb{R}$ eine Funktion und $x_0 \in U$.

1. Der Funktionswert $f(x_0)$ heißt **globales Maximum** (oder auch nur: **Maximum**) der Funktion f , wenn $f(x) \leq f(x_0) \quad \forall x \in U$. In diesem Fall heißt x_0 **Maximalstelle** von f .
2. Der Funktionswert $f(x_0)$ heißt **lokales Maximum** der Funktion f , wenn es ein offenes Intervall $]x_0 - \epsilon, x_0 + \epsilon[$ gibt mit $f(x) \leq f(x_0) \quad \forall x \in U \cap]x_0 - \epsilon, x_0 + \epsilon[$. In diesem Fall heißt x_0 **lokale Maximalstelle** von f .

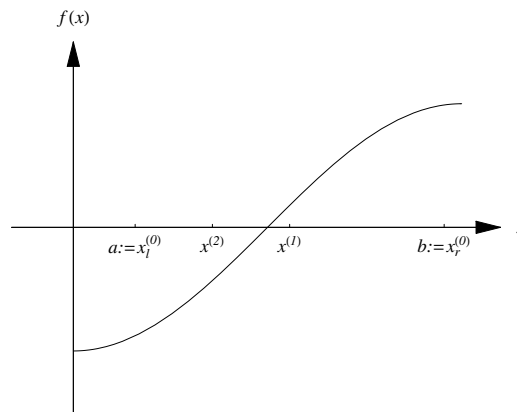


Abbildung 4.9: Intervallschachtelung

3. Ein (lokale oder globale) Maximalstelle heißt **isoliert**, wenn die Ungleichung $f(x) \leq f(x_0)$ in der jeweiligen Definition durch die strikte Ungleichung $f(x) < f(x_0)$ für $x \neq x_0$ ersetzt werden kann.
4. Analog sind globale und lokale **Minima** und (isolierte) globale und lokale **Minimalstellen** definiert.

Bemerkung 4.5.15. Jede globale Extremalstelle ist auch eine lokale. Die Umkehrung gilt aber nicht. Die in Abbildung 4.10 dargestellte Funktion besitzt ein globales Maximum in x_0 . Die in Abbildung 4.11 dargestellte Funktion hat in a ein lokales Minimum, in b ein lokales Maximum, in c ein globales Minimum und in d ein globales Maximum.

Satz 4.5.16 (Extrema einer stetigen Funktion auf kompakten Intervallen).

Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig. Dann hat f ein Maximum, d.h. $\exists x_0 \in [a, b]$ mit der Eigenschaft, dass

$$\forall x \in [a, b] \quad f(x_0) \geq f(x).$$

Ebenso nimmt f sein Minimum an.

Beweisidee:

1. f ist beschränkt: Angenommen, f sei unbeschränkt. Dann existiert eine Folge $(x^{(n)})_{n \in \mathbb{N}}$ mit $f(x^{(n)}) > n$.
Satz von Bolzano Weierstraß $\Rightarrow \exists$ eine konvergente Teilfolge $(x^{(n_k)})_{k \in \mathbb{N}}$ mit $\lim_{k \rightarrow \infty} x^{n_k} = \bar{x} \in [a, b]$. Wegen der Stetigkeit von f gilt dann aber $\lim_{k \rightarrow \infty} f(x^{n_k}) = f(\bar{x})$, was im gewünschten Widerspruch zur Unbeschränktheit von $(f(x^{(n)}))_{n \in \mathbb{N}}$ steht.
2. Das Supremum wird angenommen. Der Beweis dafür erfolgt auch mit dem Satz von Bolzano Weierstraß. □

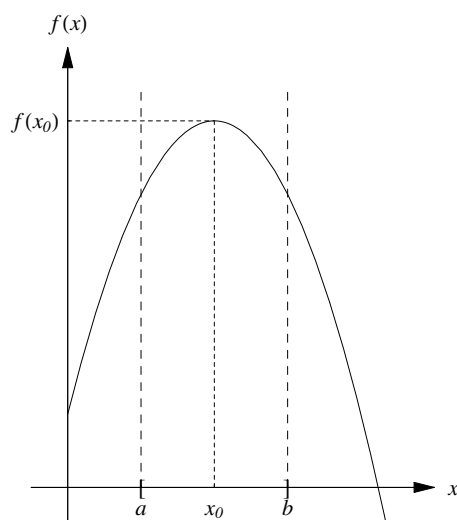


Abbildung 4.10: Die Funktion besitzt bei x_0 ein Maximum.

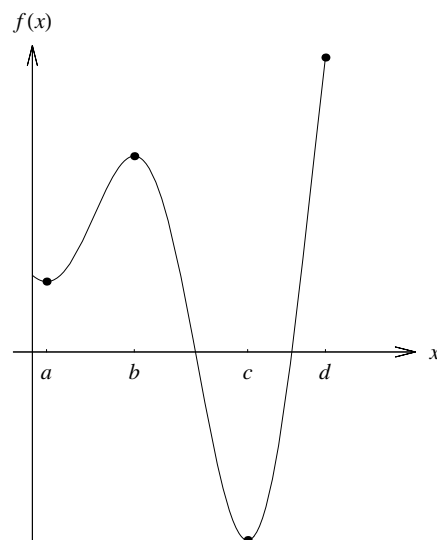


Abbildung 4.11: Die Funktion besitzt ein globales Maximum und Minimum und zusätzlich ein lokales Maximum und Minimum

Bemerkung 4.5.17.

1. Es folgt aus Satz 4.5.16 und Satz 4.5.12.2, dass kompakte Intervalle auf ebensolche surjektiv abgebildet werden.

$$f([a, b]) = [\min_{x \in [a, b]} f(x), \max_{x \in [a, b]} f(x)].$$

2. In Satz 4.5.16 ist die Beschränktheit des Intervalls $[a, b]$ notwendig für die allgemeine Schlußfolgerung:

Gegenbeispiel: $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = x$ (Bild unbeschränkt);
oder $f(x) = \arctan x$ (Bild nicht abgeschlossen).

3. Ebenso ist die Abgeschlossenheit des Intervalls notwendig. Gegenbeispiel: $f : [0, 1[\rightarrow \mathbb{R}$, $f(x) = x$. Die Funktion f nimmt ihr Supremum 1 nicht an.

Satz 4.5.18 (Inverse einer stetigen Funktion).

Seien $U \subset \mathbb{R}$ ein Intervall und $f : U \rightarrow \mathbb{R}$ eine stetige injektive Funktion. Dann gilt:

1. f ist entweder streng monoton steigend oder streng monoton fallend.
2. Sei $V := f(U)$. $f : U \rightarrow V$ ist bijektiv. Die Inverse $f^{-1} : V \rightarrow U$ ist stetig.

4.6 Differenzierbarkeit

Zur Motivation des Ableitungsbegriffes betrachten wir ein physikalische Beispiel und eine geometrische Fragestellung.

1. Durch Funktionen werden z.B. Bahnen von physikalischen Teilchen beschrieben, z.B. im eindimensionalen Raum:

$$\begin{aligned} f : [0, T] &\rightarrow \mathbb{R}, \\ t &\mapsto f(t). \end{aligned}$$

Dabei ist $f(t)$ die Position des Teilchens zur Zeit t . Man möchte auch eine Geschwindigkeit und eine Beschleunigung definieren. Diese Größen werden z.B. in der Newtonschen Mechanik benötigt.

2. Man möchte oft komplizierte Abbildungen durch einfache (affin-lineare) ersetzen, da man über diese mehr und leichter Aussagen machen oder Berechnungen anstellen kann. Die Sekante wird durch den Punkt x_0 und einen weiteren Punkt $x \neq x_0$ gebildet. Jetzt betrachtet man $x \rightarrow x_0 \Leftrightarrow h \rightarrow 0$, wobei $h := x - x_0$. Wie bei der Stetigkeit sollte die „Grenzgerade“ nicht von der Folge $x^{(n)} \rightarrow x_0$ abhängen.

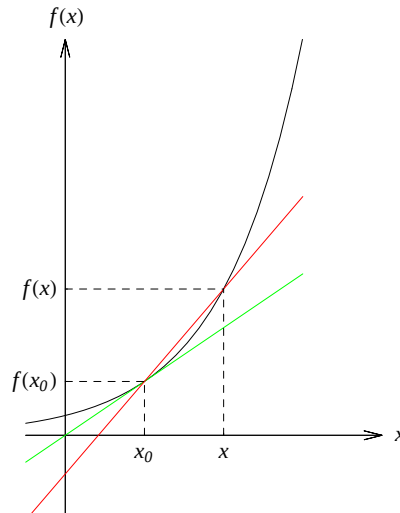


Abbildung 4.12: Die Tangente an der Stelle x_0 wird durch die Sekante angenähert

Definition 4.6.1 (Differenzierbarkeit, Ableitung).

1. Sei $U = (a, b)$ ein offenes Intervall und $x_0 \in U$. Eine Funktion $f : U \rightarrow \mathbb{R}$ heißt **differenzierbar** (genauer: **einmal differenzierbar**) in x_0 , wenn für jede Folge $(x^{(n)})_{n \in \mathbb{N}}$ mit $x^{(n)} \in U \setminus \{x_0\}$ und $\lim_{n \rightarrow \infty} x^{(n)} = x_0$ die Folge der Differenzenquotienten $\frac{f(x^{(n)}) - f(x_0)}{x^{(n)} - x_0}$ konvergiert.

Dann bezeichnen wir den Grenzwert mit

$$\begin{aligned} f'(x_0) &:= \lim_{\substack{x \rightarrow x_0 \\ x \neq x_0}} \frac{f(x) - f(x_0)}{x - x_0} \\ &= \lim_{\substack{h \rightarrow 0 \\ h \neq 0}} \frac{f(x_0 + h) - f(x_0)}{h}. \end{aligned}$$

Die Zahl $f'(x_0)$ ist die **die erste Ableitung von f an der Stelle x_0** .

2. Die Funktion f heißt **(einmal) differenzierbar auf U** , wenn sie in jedem Punkt $x_0 \in U$ (einmal) differenzierbar ist. In diesem Fall erhalten wir eine Funktion

$$f' : U \rightarrow \mathbb{R},$$

die **erste Ableitung** von f .

3. Wenn f auf U differenzierbar und die Ableitung $f' : U \rightarrow \mathbb{R}$ stetig ist, dann wird f als **einmal stetig differenzierbar** bezeichnet.

Einführung der Differentialschreibweise (nach Leibniz)

$$\begin{aligned} f'(x) = \frac{df}{dx} &:= \lim_{\Delta x \rightarrow 0} \frac{f(x + \Delta x) - f(x)}{\Delta x} \\ &\text{„}\frac{df}{dx} = \lim_{\Delta x \rightarrow 0} \frac{\Delta f}{\Delta x}\text{“} \end{aligned} \quad (4.20)$$

Interpretation von df , dx als „beliebig“ kleine Größen (Differenziale)

Streng mathematisch ist diese Interpretation problematisch, aber sehr praktisch in der Anwendung. Man „rechnet“ dort gern mit Differentialen:

$$\frac{df}{dx} = f'(x) \rightarrow df = f'(x)dx \quad (4.21)$$

(daher steht unter dem Integral immer $f(x)dx$, s. später)

Wichtige Anwendungen der Leibniz-Notation

- Thermodynamik in Physik und Chemie
- Theoretische Physik

Definition 4.6.2 (höhere Ableitungen).

1. Falls f' differenzierbar ist, dann heißt $(f')' = f''$ die **zweite Ableitung** von f . Analog definiert man die **n -te Ableitung**, vorausgesetzt, dass f hinreichend oft differenzierbar ist. Wir bezeichnen die n -te Ableitung mit $f^{(n)}$.
2. Falls $f^{(n)}$ stetig ist, wird f als **n -mal stetig differenzierbar** bezeichnet. Der Raum der n -mal stetig differenzierbaren Funktion wird mit $\mathcal{C}^n(U, \mathbb{R})$ oder auch $\mathcal{C}^n(U)$ bezeichnet.
3. Falls für jedes n , die Funktion f n -mal stetig differenzierbar ist, so wird f als **beliebig oft differenzierbar** oder auch als **glatt** bezeichnet. Der Raum der glatten Funktion ist $\mathcal{C}^\infty(U)$ oder auch $\mathcal{C}^\infty(U, \mathbb{R})$.

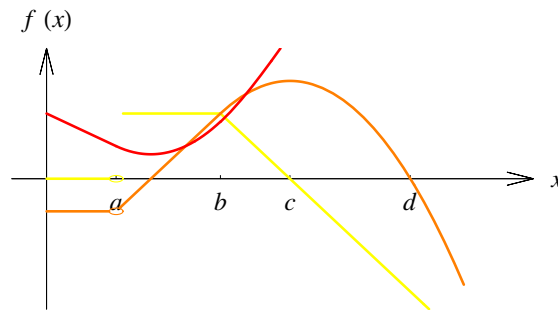


Abbildung 4.13: Eine Funktion und ihre erste und zweite Ableitung.

Bemerkung 4.6.3. $\mathcal{C}^0(U)$ ist der Raum der stetigen Funktionen.

Beispiel 4.6.4 (Ableitung einiger wichtiger Funktionen).

1. $f(x) = c$ ist eine konstante Funktion, $f^{(n)}(x) = 0$ ist glatt für $n \geq 1$.
2. $f(x) = \lambda \cdot x$, $\lambda \in \mathbb{R}$ ist glatt, $f'(x) = \lambda$, $f^{(n)} = 0$ für $n \geq 2$.
3. $f(x) = x^2$ ist glatt.

Berechnung der ersten Ableitung bei x_0 :

$$\begin{aligned} \frac{(x_0 + h)^2 - x_0^2}{h} &= \frac{x_0^2 + 2x_0h + h^2 - x_0^2}{h} \\ &= 2x_0 + h. \end{aligned}$$

Also $\lim_{h \rightarrow 0} \frac{f(x_0+h) - f(x_0)}{h} = 2x_0$, d.h. $f'(x) = 2x$.

4. $f(x) = e^x$ ist glatt. $f'(x) = e^x$ und $f^{(n)}(x) = e^x$.
5. $f(x) = \cos(x)$ ist glatt. $f'(x) = -\sin(x)$.
6. $f(x) = \sin(x)$ ist glatt. $f'(x) = \cos(x)$.
7. $f(x) = |x|$ ist glatt auf $\mathbb{R} \setminus \{0\}$, aber *nicht* differenzierbar in 0, siehe Abbildung 4.14.

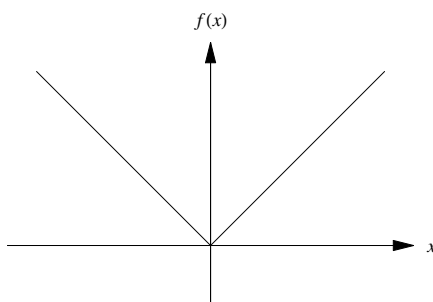
Satz 4.6.5 (Differenzierbarkeit impliziert Stetigkeit).

Sei $f : U \rightarrow \mathbb{R}$ in $x_0 \in U$ differenzierbar. Dann ist f in x_0 stetig.

Beweisidee: Aus der Konvergenz von $\frac{f(x) - f(x_0)}{x - x_0}$ für $x \rightarrow x_0$ folgt insbesondere die Konvergenz von $f(x) - f(x_0)$ gegen 0. $\square$

Bemerkung 4.6.6.

1. Nach Satz (4.6.5) ist jede differenzierbare Funktion auch stetig. Die Umkehrung gilt nicht (siehe z.B. Beispiel (4.6.4.7)).
Es gibt sogar stetige Funktionen, die in keinem Punkt differenzierbar sind. Ein Beispiel sind die „typischen“ Pfade der eindimensionalen Brownschen Bewegung.

Abbildung 4.14: Die Betragsfunktion $f(x) = |x|$.

2. **(Beispiel einer differenzierbaren Funktion, deren Ableitung nicht stetig ist)** Aus der einmaligen Differenzierbarkeit folgt nicht die stetige Differenzierbarkeit.

Gegenbeispiel (Vergleich Abbildung 4.15):

$$f(x) = \begin{cases} x^2 \cdot \cos \frac{1}{x} & \text{für } x \neq 0, \\ 0 & \text{für } x = 0. \end{cases}$$

Es gilt $f'(0) = 0$, aber „ $\lim_{x \searrow 0} f'(x)$ “ existiert nicht. Um dies zu sehen, berechnen wir

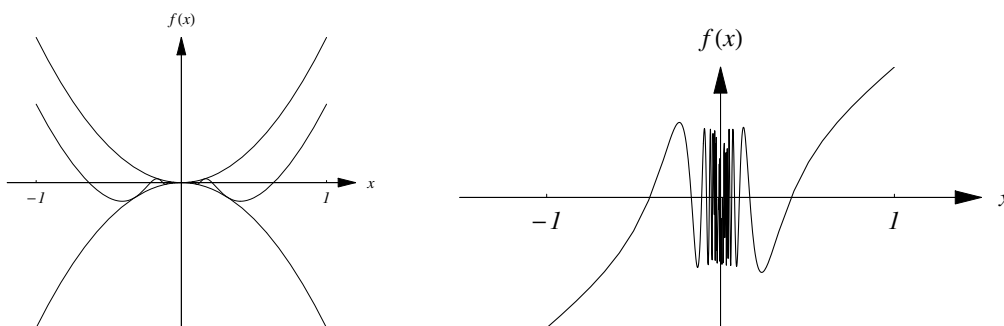


Abbildung 4.15: Graph der Funktion $f(x) = x^2 \cdot \cos \frac{1}{x}$ mit einhüllenden Parabeln (links), und ihrer Ableitung (rechts).

$f'(0)$ durch Grenzwertbildung des Differenzenquotientens und $f'(x)$ für $x \neq 0$ mit Hilfe von Produkt- und Kettenregel.

Sei $x = 0$. Wir erhalten für $h \neq 0$ unter Verwendung der Ungleichung $\cos \frac{1}{h} \leq 1$ die Abschätzung

$$\begin{aligned} \left| \frac{f(h) - f(0)}{h} \right| &= \left| \frac{1}{h} \cdot h^2 \cos \left(\frac{1}{h} \right) \right| \\ &\leq |h|, \end{aligned}$$

und somit

$$f'(0) = \lim_{\substack{h \rightarrow 0 \\ h \neq 0}} \frac{f(h) - f(0)}{h} = 0.$$

Für $x \neq 0$ gilt

$$f'(x) = 2x \cos\left(\frac{1}{x}\right) + \sin\left(\frac{1}{x}\right). \quad (4.22)$$

Die Funktion f ist also überall einmal differenzierbar und hat die Ableitung

$$f'(x) = \begin{cases} 2x \cos\left(\frac{1}{x}\right) + \sin\left(\frac{1}{x}\right) & \text{für } x \neq 0, \\ 0 & \text{für } x = 0. \end{cases}$$

Aus (4.22) erkennen wir aber auch, daß der fragliche Grenzwert „ $\lim_{x \rightarrow 0} f'(x)$ “ nicht existiert. Während der erste Summand $2x \cos\left(\frac{1}{x}\right)$ gegen 0 konvergiert, oszilliert der zweite zwischen -1 und 1 : Für die Nullfolgen $(x_n)_n$ mit $\frac{1}{x_n} = \frac{\pi}{2} + 2\pi n$ und $(y_n)_n$ mit $\frac{1}{y_n} = \frac{3\pi}{2} + 2\pi n$ gilt nämlich

$$\begin{aligned} \sin\left(\frac{1}{x_n}\right) &= 1, \\ \sin\left(\frac{1}{y_n}\right) &= -1. \end{aligned}$$

Satz 4.6.7 (Produkt- und Quotientenregel).

Seien $f, g : U \rightarrow \mathbb{R}$ (n -mal stetig) differenzierbar. Dann sind folgende Funktionen (n -mal stetig) differenzierbar:

1. $f + g$ mit

$$(f + g)'(x) = f'(x) + g'(x),$$

2. $f \cdot g$ mit

$$(f \cdot g)'(x) = f'(x) \cdot g(x) + f(x) \cdot g'(x) \quad \textbf{(Produktregel)},$$

3. (falls zusätzlich $f(x) \neq 0$ gilt) $\frac{1}{f}$ mit

$$\left(\frac{1}{f}\right)'(x) = \frac{-f'(x)}{(f(x))^2},$$

4. (falls zusätzlich $f(x) \neq 0$ gilt) $\frac{g}{f}$ mit

$$\left(\frac{g}{f}\right)'(x) = \frac{g'(x) \cdot f(x) - g(x) \cdot f'(x)}{(f(x))^2} \quad \textbf{(Quotientenregel)}.$$

Beispiel 4.6.8 (Anwendung von Produkt- und Quotientenregel).

1. (zur Produktregel)

$$\begin{aligned} f(x) &= e^x \cdot \sin x, \\ f'(x) &= e^x \sin x + e^x \cos x. \end{aligned}$$

2. (zur Quotientenregel)

$$\begin{aligned} f(x) &= \frac{x^2}{x^3 + 1}, \\ f'(x) &= \frac{2x \cdot (x^3 + 1) - x^2 \cdot 3x^2}{(x^3 + 1)^2} \\ &= \frac{-x^4 + 2x}{(x^3 + 1)^2}. \end{aligned}$$

Satz 4.6.9 (Kettenregel).

Seien $g : U \rightarrow \mathbb{R}$, $f : V \rightarrow \mathbb{R}$ n -mal stetig differenzierbar und $g(U) \subset V$. Dann ist $f \circ g : U \rightarrow \mathbb{R}$ auch n -mal stetig differenzierbar und

$$(f \circ g)'(x) = f'(g(x)) \cdot g'(x).$$

Beispiel 4.6.10 (zur Kettenregel).

1.

$$\begin{aligned} f(x) &= e^{\lambda x}, \\ f'(x) &= e^{\lambda x} \cdot \lambda. \end{aligned}$$

2.

$$\begin{aligned} f(x) &= e^{-x^2}, \\ f'(x) &= e^{-x^2} \cdot (-2x). \end{aligned}$$

3.

$$\begin{aligned} f(x) &= \sin(\cos x), \\ f'(x) &= \cos(\cos x) \cdot (-\sin x). \end{aligned}$$

Satz 4.6.11 (Differenzierbarkeit der Inversen Funktion).

Sei $f :]x_1, x_2[\rightarrow]y_1, y_2[$ n -mal stetig differenzierbar und umkehrbar, d.h. $f^{-1} :]y_1, y_2[\rightarrow]x_1, x_2[$ existiere. Desweiteren seien $x \in]x_1, x_2[$ und $f'(x) \neq 0$.

Dann ist f^{-1} an der Stelle $y = f(x)$ n -mal stetig differenzierbar und es gilt (siehe die Abbildungen 4.17 und 4.16):

$$(f^{-1})'(y) = \frac{1}{f'(x)}, \quad \text{wobei } x = f^{-1}(y).$$

□

Bemerkung 4.6.12.

1. In Satz 4.7.7 (s.u.) wird ein „handhabbares“ hinreichendes Kriterium für die (lokale) Umkehrbarkeit von differenzierbaren Funktionen angegeben.
2. Man kann sich die Formel für die Ableitung der Inversen leicht merken. Es gilt nämlich:

$$f^{-1} \circ f(x) = \text{id}(x) = x.$$

Ableiten auf beiden Seiten führt zu

$$\begin{aligned} (f^{-1})'(f(x)) \cdot f'(x) &= 1 \\ \Leftrightarrow (f^{-1})'(f(x)) &= \frac{1}{f'(x)} \end{aligned}$$

oder, äquivalent dazu:

$$(f^{-1})'(y) = \frac{1}{f'(f^{-1}(y))}.$$

Das ist aber *kein* Beweis von Satz (4.6.11). Die Umformungen sind erst gerechtfertigt, wenn Differenzierbarkeit (Voraussetzung für die Kettenregel) nachgewiesen ist.

Beispiel 4.6.13 (für Umkehrfunktionen).

1. (**Exponentialfunktion und Logarithmus**) $f : \mathbb{R} \rightarrow \mathbb{R}^{>0}$ (Wertebereich ist $\mathbb{R}^{>0} = \{y \in \mathbb{R} : y > 0\}$)

$$\begin{aligned} f(x) &= e^x = \exp(x), \\ f'(x) &= e^x = f(x). \end{aligned}$$

Die Funktion f ist streng monoton steigend. Also existiert eine Umkehrabbildung, der natürliche Logarithmus:

$$\begin{aligned} f^{-1} &= \ln : \mathbb{R}^{>0} \rightarrow \mathbb{R}, \\ y &\mapsto \ln y. \end{aligned}$$

Satz (4.6.11) liefert:

$$\begin{aligned} (f^{-1})'(y) &= \frac{1}{e^x} \\ &= \frac{1}{e^{\ln y}} \\ &= \frac{1}{y}. \end{aligned}$$

Aus den *Funktionalgleichungen* für die Exponentialfunktion:

$$\begin{aligned} e^{x_1+x_2} &= e^{x_1} \cdot e^{x_2} \quad \forall x_1, x_2, r \in \mathbb{R}, \\ (e^{x_1})^r &= e^{rx_1}, \end{aligned}$$

können wir die für den Logarithmus herleiten. Es gilt

$$\begin{aligned} \exp(\ln y_1 + \ln y_2) &= \exp(\ln y_1) \cdot \exp(\ln y_2) \\ &= y_1 \cdot y_2 \\ &= \exp(\ln(y_1 \cdot y_2)). \end{aligned}$$

Aus der Injektivität von $\exp$ folgt:

$$\ln y_1 + \ln y_2 = \ln(y_1 \cdot y_2) \quad \forall y_1, y_2 > 0.$$

Ebenso zeigt man:

$$\ln(y^r) = r \cdot \ln y \quad \forall y, r > 0.$$

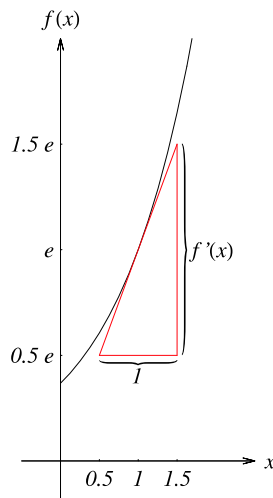


Abbildung 4.16: Die Ableitung entspricht der Steigung $\frac{f'(x)}{1}$ einer Tangente.

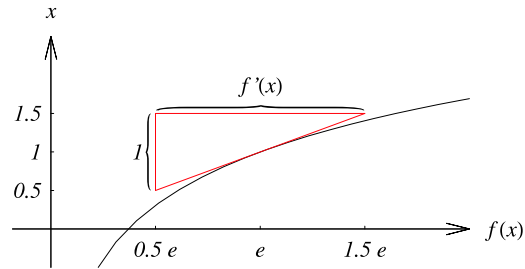


Abbildung 4.17: Die Ableitung der Umkehrfunktion entspricht der Steigung $\frac{1}{f'(x)}$ der umgekehrten Tangente.

Wegen $y > 0$ gilt nämlich

$$\begin{aligned} y = e^x &\Leftrightarrow \ln y = x \\ &\Rightarrow \ln(y^r) = \ln((e^x)^r) = \ln(e^{rx}) = r \cdot x = r \cdot \ln y. \end{aligned}$$

2. (Funktionen x^r) Sei $0 \neq r$ fest gewählt und $f : \mathbb{R}^{>0} \rightarrow \mathbb{R}^{>0}$.

$$\begin{aligned} f(x) &= x^r \\ &= \exp(\ln(x^r)) \\ &= \exp(r \cdot \ln x). \end{aligned}$$

Aus der Kettenregel folgt:

$$\begin{aligned} f'(x) &= \exp(r \cdot \ln x) \cdot r \cdot \frac{1}{x} \\ &= x^r \cdot r \cdot \frac{1}{x} \\ &= r \cdot x^{r-1}. \end{aligned}$$

(Im Fall von $r = 1$, ist $x^{r-1} = 0$ definiert.)

Insbesondere gilt für $r = \frac{1}{2}$:

$$\begin{aligned} f(x) &= \sqrt{x}, \\ f'(x) &= \frac{1}{2\sqrt{x}}. \end{aligned}$$

Die Wurzelfunktion ist also auf $R > 0$ differenzierbar. An der Stelle Null ist die Ableitung aber singulär:

$$\lim_{x \searrow 0} \frac{1}{2\sqrt{x}} = +\infty.$$

4.7 Der Mittelwertsatz

Oft interessiert man sich für Maxima und Minima einer Funktion, z.B. wenn diese einen Gewinn in Abhängigkeit von variablen Parametern darstellt. Des Weiteren können viele Naturgesetze (Modelle der Natur) als Variationsprinzip formuliert werden: „Das Licht nimmt den optisch kürzesten Weg“ (vgl. Bemerkung 4.9.6), Variationsprinzipien für die *Wirkung* („Energie mal Zeit“), z.B. in der klassischen Mechanik (nach Lagrange und anderen).

Wie findet man z.B. geeignete Kandidaten für eine Maximalstelle (Minimalstelle) einer differenzierbaren Funktion?

Satz 4.7.1. (Notwendige Bedingung für ein Maximum (Minimum) im Inneren)

Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig und differenzierbar in $x_0 \in]a, b[$. Desweiteren habe f ein (lokales) Maximum (Minimum) in x_0 , d.h. $\exists \epsilon > 0$ mit der Eigenschaft $]x_0 - \epsilon, x_0 + \epsilon[\subset]a, b[$ und

$$\forall x \in]x_0 - \epsilon, x_0 + \epsilon[\quad f(x_0) \geq f(x) \quad (\text{bzw. } f(x_0) \leq f(x)).$$

Dann gilt $f'(x_0) = 0$.

Beweis: Sei x_0 lokale Maximalstelle und ϵ wie in der Voraussetzung beschrieben. Dann gilt für $x \in]x_0 - \epsilon, x_0[$

$$\frac{f(x) - f(x_0)}{x - x_0} \geq 0,$$

also

$$f'(x_0) = \lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0} \geq 0$$

Ebenso zeigt man, indem man $x \in]x_0, x_0 + \epsilon[$ betrachtet, dass: $f'(x_0) \leq 0$, also $f'(x_0) = 0$. $\square$

Bemerkung 4.7.2. An (lokalen) Maximalstellen *am Rand* eines zumindest einseitig abgeschlossenen Intervalls $[a, b]$ (oder auch z.B. $[a, b[$) muß die Ableitung nicht notwendig verschwinden.

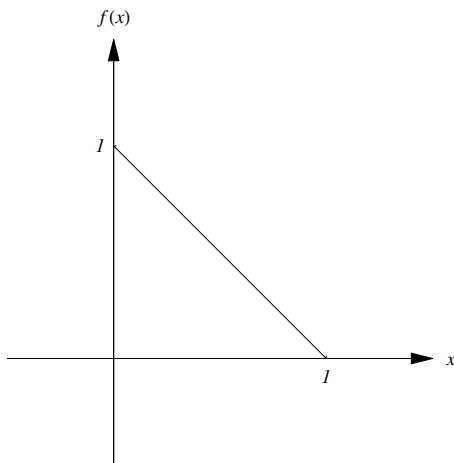
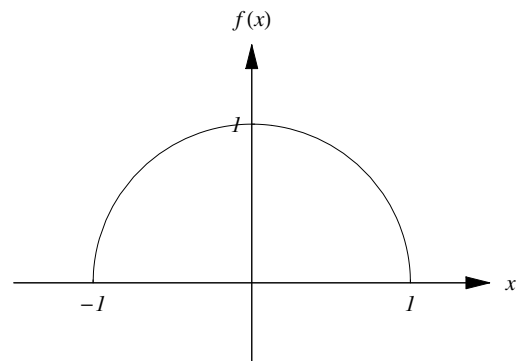
Beispiel: (vergleiche Abbildung 4.18)

$$\begin{aligned} f : [0, 1] &\rightarrow \mathbb{R}, \\ x &\mapsto 1 - x. \end{aligned}$$

Die Funktion f ist an der Stelle 0 maximal aber $f'(0) = -1$. Dabei ist $f'(0)$ als Limes der (einseitigen) Differenzenquotienten

$$\lim_{x \searrow 0} \frac{f(x) - f(0)}{x - 0} =: f'(0)$$

definiert.

Abbildung 4.18: Graph von $1 - x$ Abbildung 4.19: Graph von $\sqrt{1 - x^2}$ (Kreisbogen)**Satz 4.7.3** (Satz von Rolle).

Seien $f \in \mathcal{C}^0([a, b])$ und differenzierbar auf $]a, b[$ und $f(a) = f(b)$. Dann existiert ein $\xi \in]a, b[$ mit $f'(\xi) = 0$.

Beweis: 1. Fall: Sei f konstant auf $[a, b]$. Dann erfüllt offensichtlich jedes $\xi \in]a, b[$ die Bedingung $f'(\xi) = 0$.

2. Fall: Sei f nicht konstant auf $]a, b[$, d.h. es gibt ein $x \in]a, b[$ mit $f(x) \neq f(a)$. Sei ohne Einschränkung der Allgemeinheit $f(x) > f(a)$. Dann hat f nach Satz 4.5.16 ein Maximum $\epsilon]a, b[$ und nach Satz 4.7.1 gilt $f'(\xi) = 0$. $\square$

Beispiel 4.7.4.

1. $f : [-1, 1] \rightarrow \mathbb{R}$, $f(x) = \sqrt{1 - x^2}$ (siehe Abbildung 4.19) f ist stetig differenzierbar auf $] - 1, 1[$ und stetig auf $[-1, 1]$. Aber f ist nicht (einseitig) differenzierbar an den Stellen $-1, 1$. Desweiteren gilt $f(-1) = f(1) = 0$.

Nach dem Satz von Rolle existiert ein $\xi \in] - 1, 1[$ mit $f'(\xi) = 0$.

Bei diesem Beispiel ist ξ eindeutig und bekannt, nämlich $\xi = 0$.

2. $f : [0, \pi] \rightarrow \mathbb{R}$, $f(x) = e^{x^2} \cdot \sin x$. Es gilt $f(0) = f(\pi) = 0$.

$$\begin{aligned} f'(x) &= e^{x^2} \cdot 2x \sin x + e^{x^2} \cdot (-\cos x) \\ &= e^{x^2} \cdot [2x \sin x - \cos x] \stackrel{!}{=} 0 \end{aligned}$$

$$\Leftrightarrow 2x \sin x = \cos x$$

$$\Leftrightarrow 2x = \frac{\cos x}{\sin x} = \cot x.$$

Die Existenz eines $\xi \in]0, \pi[$ mit $2\xi = \cot \xi$ ist nach Satz 4.7.3 gewährleistet, aber man muß die Gleichung nicht unbedingt explizit lösen können.

Es gibt z.B. Polynome 5. Grades ($\Rightarrow$ mindestens eine reelle Nullstelle), deren Nullstellen man nicht „explizit“ darstellen kann.

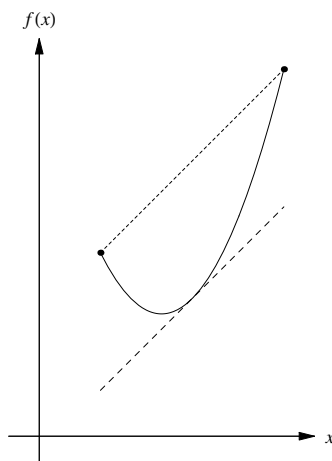


Abbildung 4.20: Die Funktion nimmt mindestens einmal die Steigung der Sekante an.

Satz 4.7.5 (Mittelwertsatz).

Sei $f \in \mathcal{C}^0([a, b], \mathbb{R})$ und f differenzierbar in $]a, b[$. Dann gibt es ein $\xi \in]a, b[$ mit

$$f'(\xi) = \underbrace{\frac{f(b) - f(a)}{b - a}}_{\text{Steigung der Sekante, siehe Abbildung 4.20}}$$

Beweis: Wende den Satz von Rolle (4.7.3) auf die Hilfsfunktion

$$\begin{aligned} g : [a, b] &\rightarrow \mathbb{R} \\ g(x) &= f(x) - \frac{x-b}{a-b}f(a) - \frac{x-a}{b-a}f(b) \end{aligned}$$

an. Es gilt

$$\begin{aligned} g(a) &= f(a) - \frac{a-b}{a-b}f(a) - 0 \cdot f(b), \\ &= 0, \\ g(b) &= 0, \\ 0 &= g'(\xi) = f'(\xi) - \frac{1}{a-b}f(a) - \frac{1}{b-a}f(b) \\ &= f'(\xi) - \frac{f(b) - f(a)}{b-a}. \end{aligned}$$

□

Bemerkung 4.7.6. Bemerkung Der Mittelwertsatz garantiert die Existenz eines solchen ξ , sagt aber nicht, ob ξ eindeutig bestimmt ist, oder wie man es findet.

Satz 4.7.7 (Monotone und konstante Funktionen).

Sei $f :]a, b[\rightarrow \mathbb{R}$ differenzierbar.

1. Falls $f'(x) \geq 0 \quad \forall x \in]a, b[$ (bzw. $f'(x) < 0 \quad \forall x \in]a, b[$), dann ist f monoton steigend (bzw. monoton fallend) auf $]a, b[$.

Bei strikter Ungleichheit, also $f'(x) > 0 \quad \forall x \in]a, b[$ (bzw. $f'(x) < 0$) ist f streng monoton.

2. f ist genau dann auf $]a, b[$ konstant, wenn

$$f'(x) = 0 \quad \forall x \in]a, b[.$$

Beweis:

1. exemplarisch für $f'(x) > 0$ (der Rest von 4.7.7.1 folgt analog):

Sei $x_1 < x_2 \in]a, b[$. Zu zeigen ist $f(x_1) < f(x_2)$.

Es gibt nach dem Mittelwertsatz ein $\xi \in]x_1, x_2[$ mit

$$\begin{aligned} \frac{f(x_2) - f(x_1)}{x_2 - x_1} &= f'(\xi) > 0 \\ \Leftrightarrow f(x_2) - f(x_1) &= \underbrace{f'(\xi)}_{>0} \cdot \underbrace{x_2 - x_1}_{>0} \\ &> 0, \end{aligned}$$

was zu zeigen war.

2. Wenn $f(x) = c \quad \forall x \in]a, b[$ dann folgt $f'(x) = 0$. Ist umgekehrt $f'(x) = 0 \quad \forall x \in]a, b[$, so folgt aus (1), dass f sowohl monoton steigend als auch fallend ist. Also ist f konstant.

□

Beispiel 4.7.8 (Tangens und Arcustangens).

Die Tangensfunktion $f :]-\frac{\pi}{2}, \frac{\pi}{2}[\rightarrow \mathbb{R}$, $f(x) = \tan x = \frac{\sin x}{\cos x}$ ist nach der Quotientenregel stetig differenzierbar, sogar glatt in $D :=]-\frac{\pi}{2}, \frac{\pi}{2}[$ (siehe Abbildung 4.21), und es gilt

$$\begin{aligned} f'(x) &= \frac{\cos^2 x + \sin^2 x}{\cos^2 x} \\ &= 1 + \tan^2 x > 0. \end{aligned}$$

Nach Satz (4.7.7.1) ist f auf D streng monoton steigend. Insbesondere ist f auf D injektiv.

Wegen $\lim_{x \rightarrow \pm \frac{\pi}{2}} \tan x = \pm \infty$ ist der Wertebereich $f(D) = \mathbb{R}$. Nach Satz (4.5.18) und Satz (4.6.11) gibt es eine glatte Umkehrfunktion (siehe Abbildung 4.22).

$$f^{-1} = \arctan : \quad \mathbb{R} \rightarrow \left] -\frac{\pi}{2}, \frac{\pi}{2} \right[$$

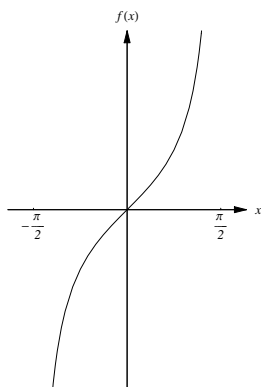


Abbildung 4.21: Die Tangensfunktion

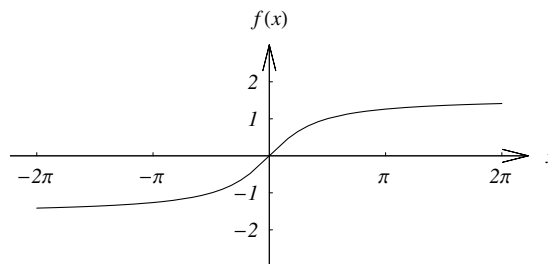


Abbildung 4.22: Die Arcustangensfunktion

mit

$$\begin{aligned}
 (f^{-1})'(y) &= \frac{1}{f'(f^{-1}(y))} \\
 &= \frac{1}{1 + [\tan(\arctan y)]^2} \\
 &= \frac{1}{1 + y^2}.
 \end{aligned}$$

Hubert Cremer [Cre79] war von dieser Kurve so fasziniert, dass er folgendes Gedicht schrieb:

Ode an die Arcustangens-Schlange

*Du schleichst seit undenklichen Zeiten
so leis und so sanft heran
Du stiegst in Ewigkeiten
kaum um ein δ an.
Nur langsam beginnst Du zu wachsen,
wie zum Beweis Deines Seins,
erreichst beim Schnittpunkt der Achsen
Deine höchste Steigung, die Eins.
Dann duckst Du Dich wieder zierlich
in stiller Bescheidenheit
und wandelst weiter manierlich
in die Unendlichkeit.*

*Hier stock ich im Lobgesange,
mir schwant, er wird mir vermiest:
Oh, Arcustangens-Schlange,
beißt du nicht doch, Du Biest ?!*

4.8 Taylorentwicklung

Sei f differenzierbar in U , $x, x_0 \in U$. Nach dem Mittelwertsatz (Satz 4.7.5) gilt

$$f(x) = \underbrace{f(x_0)}_{\text{Polynom vom Grad 0}} + \underbrace{f'(\xi) \cdot (x - x_0)}_{\text{Fehler}}.$$

Die Funktion f wird durch die konstante Funktion mit Wert $f(x_0)$ angenähert, und der Approximationsfehler ist $f(x) - f(x_0) = f'(\xi) \cdot (x - x_0)$. Wir können dies verallgemeinern, indem wir f durch Polynome höheren Grades approximieren, deren Koeffizienten durch f bestimmt sind. Wir nehmen also die Werte der Ableitung von f an der Stelle x_0 bis zum Grad n hinzu:

$$f(x_0), f'(x_0), f^{(2)}(x_0), \dots, f^{(n)}(x_0).$$

Definition 4.8.1 (Taylorpolynom und Restglied).

Sei $f : U \rightarrow \mathbb{R}$ an der Stelle $x_0 \in U$ n -mal differenzierbar.

1. Dann ist das **n -te Taylorpolynom** von f an der Stelle (Entwicklungspunkt) x_0 definiert als

$$P_n(x) = \sum_{k=0}^n \frac{f^{(k)}(x_0)}{k!} (x - x_0)^k.$$

2. Das zugehörige **Restglied** definieren wir als

$$R_n(f, x_0)(x) := f(x) - P_n(x).$$

Beispiel 4.8.2.

1. $n = 0$: $P_0(x) = f(x_0)$.
2. $n = 1$: $P_1(x) = f(x_0) + f'(x_0) \cdot (x - x_0)$.
3. $n = 2$: $P_2(x) = f(x_0) + f'(x_0) \cdot (x - x_0) + \frac{1}{2} f''(x_0) \cdot (x - x_0)^2$.

Satz 4.8.3 (Taylorsche Formel mit Restglieddarstellung nach Lagrange).

Sei $x_0 \in U$, $f \in C^{n+1}(U)$. Dann gilt

- 1.

$$f(x) = \sum_{k=0}^n \frac{f^{(k)}(x_0)}{k!} (x - x_0)^k + R_n(f, x_0)(x). \quad (4.23)$$

2. (Darstellung des Restgliedes nach Lagrange)

$$R_n(f, x_0)(x) = \frac{f^{(n+1)}(\xi)}{(n+1)!} (x - x_0)^{n+1} \quad (4.24)$$

mit einem

$$\xi \in \begin{cases}]x_0, x[& \text{falls } x > x_0, \\]x, x_0[& \text{falls } x < x_0. \end{cases}$$

Bemerkung 4.8.4 (alternative Restglieddarstellungen). Es gibt auch andere Restglieddarstellungen, z.B. nach Cauchy, Schlömilch und auch eine (leicht zu beweisende) Integraldarstellung.

Beispiel 4.8.5 (für Taylorpolynome und Taylorreihen).

1. **(Taylorreihe eines Polynoms)** Sei $f(x) = \sum_{l=0}^m a_l x^l$ ein Polynom vom Grad m . Das n -te Taylorpolynom von f an der Stelle $x_0 = 0$ ist

$$P_n(x) = \sum_{k=0}^n a_k x^k \quad \text{mit } a_k = 0 \text{ für } n > m,$$

d.h. für $n \geq m$ ist das Restglied gleich 0, da $f^{(n+1)} \equiv 0$. Insbesondere gilt für Polynome (und allgemein für absolut konvergente Potenzreihen):

$$a_k = \frac{1}{k!} P_n^{(k)}(0).$$

2. **(Taylorreihe der Exponentialfunktion)** Sei $f(x) = e^x$. Dann gilt

$$f^{(n)}(x) = e^x \quad \text{für } n \geq 1.$$

Das n -te Taylorpolynom von f für den Entwicklungspunkt $x_0 = 0$ ist wegen $e^0 = 1$ also

$$P_n(x) = \sum_{k=0}^n \frac{1}{k!} x^k.$$

Das Restglied ist $R_n(f, x_0)(x) = \frac{e^\xi}{(n+1)!} x^{n+1}$, wobei $\xi \in]0, x[$ von x und n abhängt. Für fest gewähltes x gilt in diesem Beispiel

$$\lim_{n \rightarrow \infty} R_n(f, x_0)(x) = 0$$

Also wird die Funktion $f(x) = e^x$ tatsächlich durch ihre Taylorreihe dargestellt:

$$e^x = \sum_{k=0}^{\infty} \frac{x^k}{k!} \quad (\text{Taylorreihe von } e^x)$$

Eine Illustration der ersten Partialsummen der Taylorreihe der Exponentialfunktion findet sich in Abbildung 4.23.

3. **(Beispiel einer glatten, nicht-analytischen Funktion)** Es kann auch vorkommen, dass die Taylorreihe einer Funktion f zwar konvergent ist, aber in keinem offenen Intervall um den Entwicklungspunkt gegen f konvergiert.

Gegenbeispiel:

$$f(x) = \begin{cases} 0 & \text{für } x \leq 0, \\ e^{\frac{1}{x}} & \text{für } x > 0. \end{cases}$$

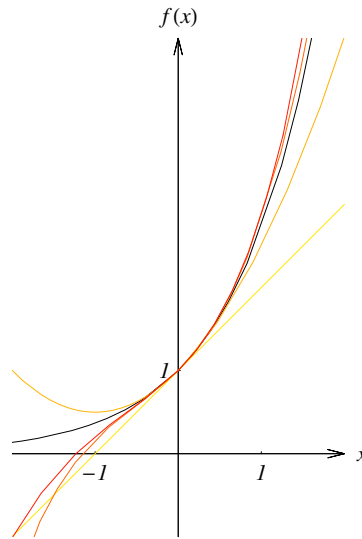


Abbildung 4.23: Die ersten Glieder der Taylorreihe der Exponentialfunktion

Es gilt $f \in \mathcal{C}^\infty(\mathbb{R}, \mathbb{R})$ und $f^{(n)}(0) = 0$. Also ist jedes Taylorpolynom und somit auch die Taylorreihe von f um den Punkt $x_0 = 0$ gleich 0. Insbesondere konvergiert die Taylorreihe auf $]0, \infty[$ nicht gegen f . Der Term $f^{(n)}(\xi)$ in der Restglieddarstellung (4.8.3.4.24) wächst „stark“ mit n und wird nicht hinreichend durch $\frac{1}{n!}$ kompensiert, d.h. die Folge der Restglieder konvergiert nicht gegen 0.

Funktionen, die sich lokal (d.h. für jeden Punkt ihres Definitionsbereiches in einer offenen Umgebung dieses Punktes) durch ihre Taylorreihe darstellen lassen, heißen *analytisch*. Die Funktion f aus diesem Beispiel ist also glatt aber nicht analytisch.

4. **(Taylorreihe der Logarithmus-Funktion)** Seien $f(x) = \ln x$ mit $x_0 = 1$. Es gilt $\ln 1 = 0$. Man kann leicht durch vollständige Induktion zeigen, dass

$$f^{(n)} = \frac{(-1)^{n+1}(n-1)!}{x^n} \quad \text{für } n \geq 1.$$

Wir können das Restglied mit Hilfe der Darstellung (4.24) abschätzen:

$$\begin{aligned} |R_n(f, x_0)(x)| &= \frac{1}{(n+1)!} \cdot n! \cdot |x - x_0|^{n+1} \\ &= \frac{|x - x_0|^{n+1}}{n+1}, \end{aligned}$$

und somit

$$\lim_{n \rightarrow \infty} |R_n(f, x_0)(x)| = 0 \quad \text{für } |x| < 1.$$

Damit ist gezeigt, dass die Taylorreihe in (4.25) mit der Funktion $\ln(1+x)$ auf dem offenen Intervall $] -1, 1[$ übereinstimmt:

$$\ln(1+x) = \sum_{n=1}^{\infty} \frac{(-1)^{n+1} x^n}{n} \quad \text{für } |x| < 1. \quad (4.25)$$

Man kann sogar zeigen, dass die Darstellung in (4.25) auch noch für $x = 1$ richtig ist. Für $x = -1$ hingegen divergiert die Reihe (harmonische Reihe), und die Funktion $\ln(1 + x)$ ist an dieser Stelle singulär.

4.9 Maxima und Minima

Mit Satz 4.5.16 hatten wir bereits ein Existenzresultat und mit Satz 4.7.1 ein notwendiges Kriterium für ein Extremum kennengelernt. (Man beachte die genauen Voraussetzungen in den jeweiligen Sätzen!) Ein Beispiel für eine Funktion, die in einem Punkt die notwendige Bedingung $f'(x) = 0$ erfüllt aber dennoch kein Extremum besitzt, ist in Abbildung 4.24(b) zu sehen. Offensichtlich reicht diese Bedingung nicht aus, um ein Extremum zu garantieren.

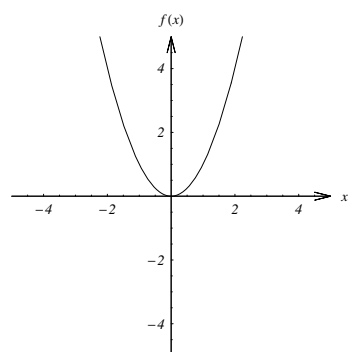
In diesem Kapitel formulieren wir *hinreichende* Kriterien für Extrema.

Satz 4.9.1 (hinreichendes Kriterium für ein Extremum).

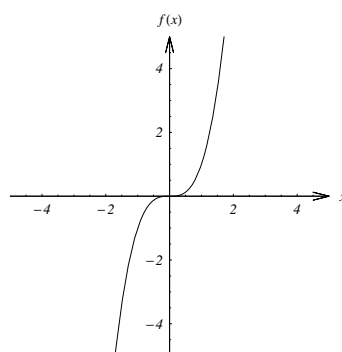
Sei $f : U \rightarrow \mathbb{R}$, $U = (a, b)$ offen in U differenzierbar (d.h. an jeder Stelle $x \in U$ differenzierbar). Im Punkt $x_0 \in U$ sei f zweimal differenzierbar und es gelte

$$\begin{aligned} f'(x_0) &= 0 \\ f''(x_0) &> 0 \quad (\text{bzw. } f''(x_0) < 0). \end{aligned}$$

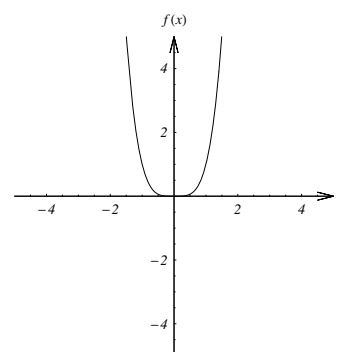
Dann ist x_0 eine isolierte lokale Minimalstelle (bzw. Maximalstelle) von f .



(a) Die Funktion $f(x) = x^2$ hat ein globales Minimum bei $x = 0$



(b) Die Funktion $f(x) = x^3$ hat eine Wendestelle bei $x = 0$



(c) Die Funktion $f(x) = x^4$ hat ein globales Minimum bei $x = 0$, aber $f''(0) = 0$.

Abbildung 4.24: Minima und Wendestellen von Funktionen $f(x) = x^n$

Beweis: Sei $f''(x_0) > 0$ (Der Fall $f''(x_0) < 0$ wird analog behandelt.)

Da

$$f''(x_0) = \lim_{x \rightarrow x_0} \frac{f'(x) - f'(x_0)}{x - x_0} > 0$$

mit $x = x_0 + h$, existiert ein $\epsilon > 0$, so dass

$$\frac{f'(x) - f'(x_0)}{x - x_0} > 0 \quad \forall x \text{ in } U_\epsilon(x_0).$$

Da $f'(x_0) = 0$ folgt

$$\begin{aligned} f'(x) &< 0 && \text{für } x_0 - \epsilon < x < x_0, \\ f'(x) &> 0 && \text{für } x_0 < x < x_0 + \epsilon. \end{aligned}$$

Nach unserem Monotoniekriterium ist also f in $[x_0 - \epsilon, x_0]$ streng monoton fallend und in $[x_0, x_0 + \epsilon]$ streng monoton steigend. $\square$

***Bemerkung 4.9.2. (Degenerierte kritische Punkte, Extrema und Wendestellen)**

1. Satz (4.9.1) gibt eine hinreichende, aber nicht notwendige Bedingung für lokale Extrema an. So hat $f(x) = x^4$ bei $x = 0$ ein isoliertes lokales Minimum, aber $f''(0) = 0$ (siehe Abbildung 4.24(c)).
2. Wir verallgemeinern die Aussage von Bemerkung 4.9.2.1. Wir sehen leicht, dass für die Funktion $f_n = x^n$ mit $(n \geq 1)$ folgendes gilt:

$$\begin{aligned} f_n^{(k)}(0) &= 0 && \text{für } 0 \leq k < n, \\ f_n^{(n)}(0) &= n! > 0. \end{aligned}$$

Falls n ungerade ist, so ist $f_n(x) < 0$ für $x < 0$ und $f_n(x) > 0$ für $x > 0$. Insbesondere hat f_n kein Extremum an der Stelle 0.

Ist n jedoch gerade, so hat f_n an der Stelle 0 ein Minimum.

3. Noch allgemeiner als in Bemerkung 4.9.2.2 betrachten wir nur ein $f \in \mathcal{C}^n(\mathbb{R})$ und ein $x_0 \in \mathbb{R}$ mit

$$\begin{aligned} f^{(k)}(x_0) &= 0 && \text{für } 1 \leq k < n, \\ f^{(n)}(x_0) &= n! \neq 0. \end{aligned}$$

Die Untersuchung von f auf Extrema oder Wendepunkte führt man mit Hilfe des n -ten Taylorpolynoms von f mit Entwicklungspunkt x_0 auf Bemerkung 4.9.2.2 zurück. Es gilt

$$P_n(x) = \frac{f^{(n)}(x_0)}{n!}(x - x_0)^n + f(x_0),$$

und „ $f(x) - f(x_0)$ “ verhält sich nahe bei x_0 so wie $P_n(x) - f(x_0)$. Insbesondere haben diese beiden Funktionen an der Stelle x_0 entweder beide ein Minimum oder ein Maximum oder den gleichen Vorzeichenwechsel.

***Definition 4.9.3 (Konvexität und Konkavität von Funktionen).**

Sei $U \subset \mathbb{R}$ ein Intervall. Eine Funktion $f : U \rightarrow \mathbb{R}$ heißt **konvex**, wenn für alle $x_1, x_2 \in U$ und alle λ mit $0 < \lambda < 1$ die Ungleichung

$$f(\lambda x_1 + (1 - \lambda)x_2) \leq \lambda f(x_1) + (1 - \lambda)f(x_2)$$

gilt (siehe Abbildung 4.25) Die Funktion f heißt **konkav**, wenn $-f$ konvex ist.

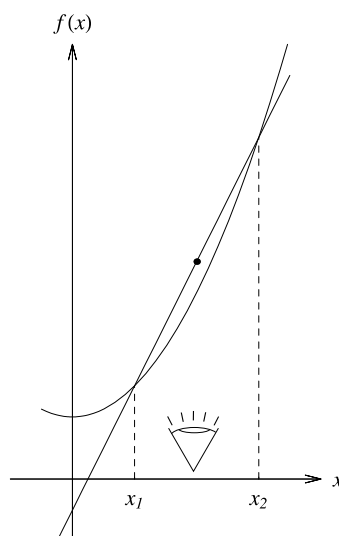


Abbildung 4.25: Der Graph einer konvexen Funktion hat einen Bauch, wenn man ihn von unten betrachtet. Ein etwas antiker Merkspruch: „Konvex ist der Bauch vom Rex.“

***Satz 4.9.4** (Konvexitätskriterium zweimal differenzierbarer Funktionen).

Sei $U \subset \mathbb{R}$ offen und $f : U \rightarrow \mathbb{R}$ eine zweimal differenzierbare Funktion. f ist genau dann konvex, falls $f''(x) \geq 0 \quad \forall x \in U$. $\square$

***Satz 4.9.5** (hinreichendes Kriterium für ein globales Extremum).

Sei $f(x)$ stetig in $U = [a, b]$ und differenzierbar in (a, b) . Hat $f(x)$ an der Stelle $x_0 \in (a, b)$ ein lokales Extremum und ist x_0 die einzige Nullstelle von f' in (a, b) , dann ist $f(x_0)$ sogar globales Extremum von $f(x)$ über $[a, b]$.

Beweis: Es ist $f(x) \neq f(x_0) \quad \forall x$ mit $a \leq x < x_0$, da sonst nach dem Satz von Rolle zwischen x und x_0 eine weitere Nullstelle der Ableitung wäre. Also ist entweder $f(x) > f(x_0)$ oder $f(x) < f(x_0) \quad \forall x$ mit $a \leq x < x_0$. Wenn $f(x_0)$ lokales Maximum ist, muß letzteres gelten und analog dazu auch $f(x) < f(x_0)$ für $x_0 < x \leq b$.

Also ist das relative Maximum zugleich globales Maximum. Der Beweis im Fall eines Minimums ist analog. $\square$

4.9.1 *Eine Optimierungsaufgabe

Ein Teilchen bewegt sich in der x, y -Ebene unterhalb der x -Achse geradlinig mit der Geschwindigkeit v_1 , oberhalb geradlinig mit der Geschwindigkeit v_2 . Auf welchem Weg kommt es am schnellsten von einem Punkt $(0, -u)$ zu einem Punkt (a, b) ?

Seien a, b, u positiv.

Frage: Wie groß ist die minimale Zeit, um von $(0, -u)$ nach (a, b) zu gelangen? Die benötigte Zeit $t(x)$ hängt nur von der Wahl von $(x, 0)$ ab! Es ist

$$t(x) = \frac{s_1}{v_1} + \frac{s_2}{v_2} = \frac{1}{v_1} \sqrt{u^2 + x^2} + \frac{1}{v_2} \sqrt{(a-x)^2 + b^2}$$

Die Funktion t ist zu minimieren. Die Formel für $t(x)$ gilt auch für negative x und $x > a$. Die Ableitung von $t(x)$ berechnen wir mit der Kettenregel

$$t'(x) = \frac{1}{v_1} \cdot \frac{x}{\sqrt{u^2 + x^2}} - \frac{1}{v_2} \cdot \frac{(a-x)}{\sqrt{(a-x)^2 + b^2}}$$

Also

$$t'(x) = \frac{1}{v_1} \cdot \frac{x}{s_1} - \frac{1}{v_2} \cdot \frac{(a-x)}{s_2}$$

Es ist $\frac{x}{s_1} = \sin \alpha$ und $\frac{(a-x)}{s_2} = \sin \beta$. Ein Kriterium für ein lokales Extremum lautet also

$$\frac{\sin \alpha}{v_1} - \frac{\sin \beta}{v_2} = 0 \quad (\text{Snellius'sches Brechungsgesetz}) \quad (4.26)$$

Gibt es genau ein x_0 , so dass (4.26) gilt? Zu berechnen wäre die zweite Ableitung. Wir können aber auch folgendermaßen argumentieren: Der Term $\sin \alpha$ wächst für $0 \leq x \leq a$ streng monoton in x , während $\sin \beta$ streng monoton fällt, also ist (4.26) nur an einer Stelle in $[0, a]$ erfüllt.

Für $x = 0$ ist $\alpha = 0$ und damit $\sin \alpha = 0$, $\sin \beta > 0$.

Für $x = a$ ist $\beta = 0$ und damit $\sin \alpha > 0$, $\sin \beta = 0$.

Also wechselt $\frac{\sin \alpha}{v_1} - \frac{\sin \beta}{v_2}$ das Vorzeichen in $[a, b]$, nach dem Zwischenwertsatz gibt es ein x_0 , so dass (4.26) erfüllt ist. ($\frac{\sin \alpha}{v_1} - \frac{\sin \beta}{v_2}$ ist stetig!)

Dieses lokale Minimum ist sogar globales Minimum:

Bemerkung 4.9.6. Es ist ein berühmtes physikalisches Prinzip, dass Licht den lokal kürzesten optischen Weg nimmt. Siehe z.B. Kapitel 26 in [FLS63].

4.10 Integration

Wir betrachten eine stetige Funktion $f : [a, b] \rightarrow \mathbb{R}$, wobei $a, b \in \mathbb{R}$ und $a < b$.

Frage: Wie groß ist der Flächeninhalt zwischen dem Abschnitt $[a, b]$ auf der x -Achse und dem Graph von f ? Zur Beantwortung dieser Frage müssen wir insbesondere einen solchen Flächeninhalt sinnvoll definieren. Das wird uns auf den Begriff des **Integrals** führen, den wir zu Beginn dieses Kapitels mathematisch exakt definieren wollen.

Wir betrachten zunächst einige einfache Beispiele.

Beispiel 4.10.1 (Integral für konstante Funktionen).

Sei f konstant und positiv, also $f(x) = c \quad \forall x \in [a, b]$ mit $c > 0$. Der fragliche Flächeninhalt ist offensichtlich der eines Rechtecks, also gleich $(b-a)c$. Wir schreiben

$$\int_a^b f(x) dx := (b-a)c. \quad (4.27)$$

Die linke Seite in (4.27) ist das **Integral** von f in den Grenzen von a bis b .

Bemerkung 4.10.2. Die Definition in (4.27) soll auch für $c < 0$ gelten. In diesem Fall ist der Flächeninhalt negativ.

Beispiel 4.10.3 (Integral für Treppenfunktionen).

Sei f ist eine Treppenfunktion, d.h. es gibt eine Zerlegung $\Delta = (x_0, \dots, x_n)$ von $[a, b]$ mit $a = x_0 < x_1 < \dots < x_n = b$, und auf jedem der offenen Teilintervalle $]x_{i-1}, x_i[$ ist die (Einschränkung) von f konstant: $f|_{]x_{i-1}, x_i[} = c_i$. Dann definieren wir das **Integral** von f in den Grenzen von a bis b als

$$\int_a^b f(x) dx := \sum_{i=1}^n (x_i - x_{i-1}) c_i. \quad (4.28)$$

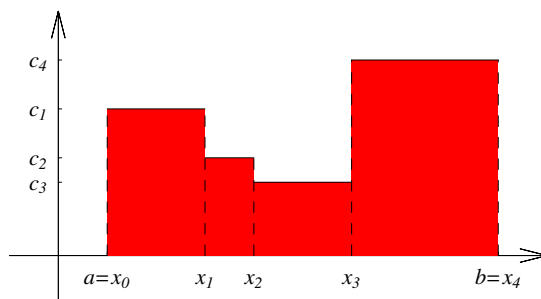


Abbildung 4.26: Das Integral einer Treppenfunktion

Satz 4.10.4 (Eigenschaften des Integrals für Treppenfunktionen).

Das in Beispiel 4.10.3 definierte Integral für Treppenfunktionen hat folgende Eigenschaften.

1. Es ist unabhängig von der Zerlegung. (Man kann ja die Funktion f mit Hilfe einer anderen (feineren) Zerlegung darstellen). Insbesondere ist das Integral als Eigenschaft der Treppenfunktion (nicht von deren spezieller Darstellung) wohldefiniert.
2. Es ist **linear** auf dem reellen Vektorraum der Treppenfunktionen auf $[a, b]$, d.h. für solche Funktionen f_1, f_2 und $\alpha \in \mathbb{R}$ gilt

$$\int_a^b (f_1 + \alpha f_2)(x) dx = \int_a^b f_1(x) dx + \alpha \int_a^b f_2(x) dx.$$

3. Es ist **monoton**: Aus der Ungleichung $f_1 \leq f_2$ (d.h. $f_1(x) \leq f_2(x) \forall x \in [a, b]$) für die Treppenfunktionen folgt die entsprechende Ungleichung für deren Integrale:

$$\int_a^b f_1(x) dx \leq \int_a^b f_2(x) dx.$$

4. Es ist **nicht-negativ**: Aus $0 \leq f$ folgt

$$0 \leq \int_a^b f(x) dx.$$

4.10.1 *Definition des Riemann-Integrals

Wir werden nun das Integral für eine allgemeinere Menge von Funktionen definieren, wobei wir einer Argumentation Riemanns folgen. Das so definierte Integral heisst mathematisch korrekt das **Riemann-Integral**, um es von anderen Definitionen des Integrals zu unterscheiden, z.B. dem sogenannten Lebesgue-Integral, die aber in dieser Vorlesung nicht ausführlich behandelt werden.

Zur Definition des Riemann-Integrals benötigen wir einige Vorbereitungen.

Definition 4.10.5 (Feinheit einer Zerlegung).

Die **Feinheit** einer **Zerlegung** $\Delta = (x_0, \dots, x_n)$ ist definiert als

$$\eta(\Delta) := \max_{1 \leq i \leq n} |x_i - x_{i-1}|.$$

Definition 4.10.6 (Ober- und Untersumme).

Sei $f : [a, b] \rightarrow \mathbb{R}$ eine beschränkte Fkt. und sei $\Delta = (x_0, \dots, x_n)$ eine Zerlegung von $[a, b]$. Dann definieren wir die **Obersumme von f bzgl. Δ** als

$$O(f, \Delta) := \sum_{i=1}^n \left(\sup_{x \in [x_{i-1}, x_i]} f(x) \right) (x_i - x_{i-1}),$$

und die **Untersumme von f bzgl. Δ** als

$$U(f, \Delta) := \sum_{i=1}^n \left(\inf_{x \in [x_{i-1}, x_i]} f(x) \right) (x_i - x_{i-1}).$$

Bemerkung 4.10.7. Die Obersumme (bzw. Untersumme) von f bzgl. einer Zerlegung Δ ist das Integral einer Treppenfunktion, die auf jedem Teilintervall $]x_{i-1}, x_i[$ konstant mit Wert $\sup_{x \in [x_{i-1}, x_i]} f(x)$ (bzw. $\inf_{x \in [x_{i-1}, x_i]} f(x)$) ist (s. Figur 4.27). (Eine solche Treppenfunktion ist bis auf die beliebige Wahl der Funktionswerte an den Stellen x_i eindeutig bestimmt und somit auch ihr Integral.)

Definition 4.10.8 (Ober- und Unterintegral).

Sei $f : [a, b] \rightarrow \mathbb{R}$ beschränkt. Wir definieren das **Oberintegral** von f als

$$\int_a^{b*} f(x) dx := \lim_{\eta(\Delta) \rightarrow 0} O(f, \Delta),$$

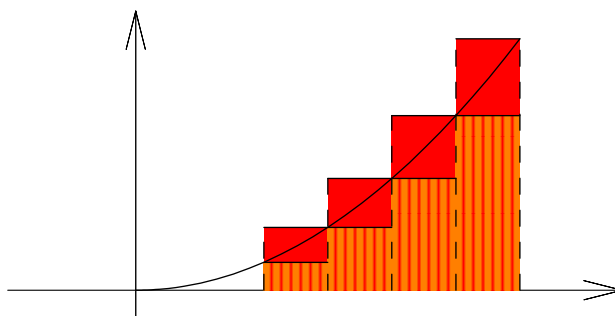


Abbildung 4.27: Ober- und Untersumme

und das **Unterintegral** von f als

$$\int_{a*}^b f(x) dx := \lim_{\eta(\delta) \rightarrow 0} U(f, \Delta).$$

Bemerkung 4.10.9. 1. Details zur Art der Grenzwertbildung in Definition 4.10.8 können z.B. in [Fora] nachgelesen werden.

2. Das Oberintegral ist größer als das Unterintegral:

$$\int_a^{b*} f(x) dx \geq \int_{a*}^b f(x) dx. \quad (4.29)$$

Definition 4.10.10 (Riemann-Integral).

Eine beschränkte Funktion $f : [a, b] \rightarrow \mathbb{R}$ heißt **Riemann-integrierbar auf dem Intervall** $[a, b]$, wenn ihre Ober- und Unterintegral gleich sind. In diesem Fall bezeichnen wir diesen Wert als **Riemann-Integral von f in den Grenzen von a bis b** :

$$\int_a^b f(x) dx := \int_a^{b*} f(x) dx.$$

Man möchte i.a. bei der Untersuchung einer gegebenen Funktion auf Integrierbarkeit natürlich nicht bei der Betrachtung von Ober- und Untersummen beginnen. Der folgende Satz garantiert die Integrierbarkeit einer großen Klasse von Funktionen.

Satz 4.10.11. (Integrierbarkeit stetiger Funktionen auf kompakten Intervallen)

1. Jede auf dem abgeschlossenen Intervall $[a, b]$ stetige Funktion f ist (auf diesem Intervall) integrierbar.
2. Jede auf dem abgeschlossenen Intervall $[a, b]$ beschränkte Funktion f mit höchstens endlich vielen Unstetigkeitsstellen ist (auf diesem Intervall) integrierbar.

Beispiel 4.10.12 (für eine nicht Riemann-integrierbare Funktion).

Wir betrachten das Beispiel

$$\begin{aligned} f : [0, 1] &\rightarrow \mathbb{R} \\ x &\mapsto \begin{cases} 1 & \text{falls } x \in \mathbb{Q} \cap [0, 1] \quad (\text{d.h. } x \text{ rational}), \\ 0 & \text{falls } x \notin \mathbb{Q} \cap [0, 1] \quad (\text{d.h. } x \text{ irrational}). \end{cases} \end{aligned}$$

Dann gilt

$$\int_{0*}^1 f(x) dx = 0 \neq 1 = \int_0^{1*} f(x) dx,$$

und somit ist die Funktion *nicht* Riemann-integrierbar.

Jetzt geben wir die Definition des Integrals für den Fall an, dass die untere Grenze nicht kleiner ist als die obere Grenze.

Definition 4.10.13. 1. Sei $f : [a, b] \rightarrow \mathbb{R}$ integrierbar. Wir definieren

$$\int_b^a f(x) dx := - \int_a^b f(x) dx.$$

2. Für eine im Punkt $a \in \mathbb{R}$ definierte Funktion f definieren wir

$$\int_a^a f(x) dx := 0.$$

Bemerkung 4.10.14. Wir werden im folgenden der Kürze halber meistens den Namen *Riemann* weglassen und nur von *Integral*, *Integrierbarkeit* usw. sprechen. Wir machen jedoch darauf aufmerksam, dass es auch andere Integraldefinitionen gibt, die in wenigen problematischen Fällen wie z.B. Beispiel 4.10.12 anders vorgehen. Für alle Funktionen, die uns in diesem Skript interessieren, reicht die Riemann-Integraldefinition jedoch aus.

Satz 4.10.15 (Eigenschaften des Integrals).

1. Seien $f : [a, b] \rightarrow \mathbb{R}$ integrierbar und $c \in]a, b[$. Dann gilt

$$\int_a^c f(x) dx + \int_c^b f(x) dx = \int_a^b f(x) dx.$$

Damit soll insbesondere gesagt sein, dass f auch auf jedem Teilintervall von $[a, b]$ integrierbar ist.

2. Das Integral ist eine **monotone** und **nicht-negative lineare Abbildung** auf dem **Vektorraum der integrierbaren Funktionen** eines Intervalls $[a, b]$. (Vgl. Satz 4.10.4.)

4.10.2 Einige Sätze zum Integral

***Satz 4.10.16** (Mittelwertsatz der Integralrechnung).

Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig. Dann existiert ein $\xi \in]a, b[$ mit

$$\int_a^b f(x) dx = (b - a)f(\xi).$$

***Satz 4.10.17** (Abschätzung des Integrals).

Sei $f : [a, b]$ integrierbar. Dann gelten die Abschätzungen

$$(b - a) \inf_{x \in [a, b]} f(x) \leq \int_a^b f(x) dx \leq (b - a) \sup_{x \in [a, b]} f(x).$$

Wir betrachten nun eine der Integrationsgrenzen als variabel.

Satz 4.10.18. (Zusammenhang zwischen Differential- und Integralrechnung)

Seien $f : [a, b] \rightarrow \mathbb{R}$ stetig und $x, a_0 \in [a, b]$. Wir definieren

$$F(x) := \int_a^x f(y) dy.$$

Dann ist $F :]a, b[\rightarrow \mathbb{R}$ differenzierbar und es gilt $F' = f$.

Beweis: Wir betrachten für festes $x \in [a, b[$ positive h , für die $x + h \leq b$ (vgl. Abbildung 4.28.) Dann ist der Differenzenquotient in (4.30) definiert. Nach Satz 4.10.16 gibt es ein (von h abhängiges) $\xi_h \in]x, x + h[$, welches folgende Gleichung erfüllt.

$$\frac{F(x + h) - F(x)}{h} = \frac{1}{h} \int_x^{x+h} f(y) dy = f(\xi_h). \quad (4.30)$$

Wegen der Stetigkeit von f gilt dann für den Grenzwert

$$\lim_{h \searrow 0} \frac{F(x + h) - F(x)}{h} = f(x).$$

Betrachtungen mit $h < 0$ oder $x = a$ oder $x = b$ sind analog dazu. □

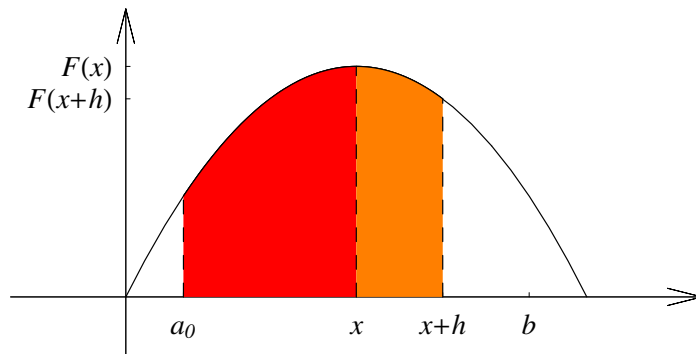
Definition 4.10.19 (Stammfunktion).

Eine differenzierbare Funktion $F : [a, b] \rightarrow \mathbb{R}$ heißt **Stammfunktion** von $f : [a, b] \rightarrow \mathbb{R}$, falls

$$F' = f. \quad (4.31)$$

Satz 4.10.20. (Eindeutigkeit der Stammfunktion bis auf eine Konstante)

Seien F und G Stammfunktionen von $f : [a, b] \rightarrow \mathbb{R}$. Dann ist die Funktion $F - G : [a, b] \rightarrow \mathbb{R}$ konstant.

Abbildung 4.28: Zuwachs der Stammfunktion über dem Intervall $[x, x + h]$

Beweis: Der Beweis folgt unmittelbar aus der Definition 4.10.19 und aus dem Mittelwertsatz der Differentialrechnung. $\square$

Aus den bisherigen Überlegungen zu Stammfunktionen folgt der folgende wichtige Satz, der eine analytische Berechnung eines Integrals auf das Auffinden einer Stammfunktion und deren Auswertung an den Integrationsgrenzen reduziert. Durch diesen Satz, Satz 4.10.18 und (4.31) ist die enge Beziehung zwischen Differential- und Integralrechnung herausgestellt.

Satz 4.10.21. (Fundamentalsatz der Differential- und Integralrechnung) Seien $f : [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion und F eine Stammfunktion von f . Dann gilt für alle $x_0, x_1 \in [a, b]$

$$\int_{x_0}^{x_1} f(x) dx = F(x_1) - F(x_0).$$

Bemerkung 4.10.22. Man verwendet oft folgende Notation:

$$F(x) \Big|_{x_0}^{x_1} := F(x_1) - F(x_0), \quad (4.32)$$

$$\int f(x) dx = F(x) + c, \quad (4.33)$$

$$\int f(x) dx \Big|_{x=g(y)} := F(g(y)). \quad (4.34)$$

Die nicht ganz saubere aber sehr praktische Notation in (4.33) bedeutet, dass F eine Stammfunktion von f ist. Die beliebig wählbare Konstante c wird oft auch weggelassen.

Die Notation auf der linken Seite von (4.34) ist so zu verstehen, dass in einer von der Variable x abhängigen Stammfunktion F von f die Substitution $x = g(y)$ vorzunehmen ist (d.h. erst integrieren, dann substituieren.)

Beispiel 4.10.23 (für Stammfunktionen).

Wir geben nun einige Beispiele von Stammfunktionen F zu Funktionen f an, die bereits aus der Differentialrechnung bekannt sind.

1. $f(x) = x^\alpha$ mit $\alpha \in \mathbb{R}$. Wir unterscheiden folgende Fälle für α .

(a) $\alpha \notin \{-1, 0\}$, $x \neq 0$. Des Weiteren setzen wir $x > 0$ voraus, falls $\alpha < 0$. Dann

$$F(x) = \frac{1}{\alpha + 1} x^{\alpha+1} + c.$$

(b) Für $\alpha = 0$ ist $f(x) = 1$. (Für $x \neq 0$ ist das klar. An der Stelle $x = 0$ haben wir f durch die stetige Fortsetzung definiert.) Dann gilt

$$F(x) = x + c.$$

(c) Für $\alpha = -1$, also $f(x) = \frac{1}{x}$, und $x \neq 0$ erhalten wir

$$F(x) = \ln |x| + c.$$

2. Für ein Polynom $f(x) = \sum_{n=0}^N a_n x^n$ gilt

$$F(x) = \sum_{n=0}^N \frac{1}{n+1} a_n x^{n+1} + c.$$

3. (a) $\int \sin x \, dx = -\cos x + c.$

(b) $\int \cos x \, dx = \sin x + c.$

4. $f(x) = e^x$, $F(x) = e^x + c.$

4.10.3 Rechenregeln zur Integration

Aufgrund der im vorangegangenen Abschnitt festgestellten Beziehung zwischen Differential- und Integralrechnung können wir aus einigen Regeln zur Ableitung von Funktionen solche über Stammfunktionen gewinnen. Die partielle Integration (Satz 4.10.24) entspricht der Produktregel und die Substitutionsregel (Satz 4.10.27) der Kettenregel.

Satz 4.10.24 (Partielle Integration).

Seien $f, g : [a, b] \rightarrow \mathbb{R}$ zwei stetig differenzierbare Funktionen. Dann gilt

$$\int_a^b f(x) \cdot g'(x) \, dx = f(x)g(x) \Big|_a^b - \int_a^b g(x)f'(x) \, dx. \quad (4.35)$$

Beweis: Wir wenden erst den Fundamentalsatz an und dann auf den Integranden die Produktregel $(f \cdot g)' = f' \cdot g + f \cdot g'$:

$$\begin{aligned} fg|_a^b &= \int_a^b (f \cdot g)'(x) \, dx \\ &= \int_a^b f'(x)g(x) \, dx + \int_a^b f(x)g'(x) \, dx. \end{aligned}$$

Durch Umformung erhalten wir (4.35). □

Bemerkung 4.10.25 (Idee der partiellen Integration).

Zur Anwendung der partiellen Integration (4.35) muss zunächst nur für einen Faktor des Integranden eine Stammfunktion gefunden werden. Es wird also nur eine Teil integriert. Dies erklärt den Namen *partielle Integration*. Von dem restlichen Faktor muss man nur die Ableitung kennen.

Beispiel 4.10.26 (zur partiellen Integration).

1. Wir suchen eine Stammfunktion zu $x e^x$. Wir beobachten, dass der Faktor x eine besonders einfache Ableitung hat. Daher nehmen wir folgende „Rollenverteilung“ vor: Wir setzen $f(x) = x$, also $f'(x) = 1$, und $g(x) = e^x$, also $g'(x) = e^x$ und erhalten

$$\begin{aligned} \int_a^b x \cdot e^x dx &= x \cdot e^x \Big|_a^b - \int_a^b e^x dx \\ &= (x \cdot e^x - e^x) \Big|_a^b. \end{aligned}$$

Mit unserer Notation (4.33) schreiben wir dies kurz als

$$\int x e^x dx = x e^x - e^x + c.$$

2. (Ergänzung des Faktors 1)

Wir möchten eine Stammfunktion von $\ln x$ für $x > 0$ berechnen. Wir kennen aber bislang nur die Ableitung dieser Funktion. Im Hinblick auf Bemerkung 4.10.25 ergänzen wir im Integranden den Faktor 1, zu dem wir natürlich eine Stammfunktion kennen, und erhalten mit $f(x) = \ln x$, $f' = \frac{1}{x}$, $g(x) = x$, $g'(x) = 1$:

$$\begin{aligned} \int \ln x dx &= \int 1 \cdot \ln x dx \\ &= x \cdot \ln x - \int x \cdot \frac{1}{x} dx \\ &= x \cdot \ln x - \int 1 dx \\ &= x \cdot \ln x - x + c. \end{aligned}$$

3. („Phoenix aus der Asche“)

In diesem Beispiel integrieren wir zweimal hintereinander partiell. Dabei wählen wir in beiden Schritten e^x als den zu integrierenden und die jeweilige trigonometrische Funktion als den abzuleitenden Faktor. (Umgekehrt ginge es hier auch.)

$$\begin{aligned} \int e^x \sin x dx &= e^x \sin x - \int e^x \cos x dx \\ &= e^x \sin x - \left[e^x \cos x + \int e^x \sin x dx \right] \\ &= e^x (\sin x - \cos x) - \int e^x \sin x dx. \end{aligned}$$

Das zu berechnende Integral ist also nach zweimaliger partieller Integration wieder aufgetaucht (daher der Name). Durch Auflösen erhalten wir

$$\int e^x \sin x \, dx = \frac{1}{2} e^x (\sin x - \cos x).$$

Satz 4.10.27 (Substitutionsregel).

Sei $g : [a, b] \rightarrow \mathbb{R}$ stetig differenzierbar, und sei f stetig auf dem Bildbereich von g . Also ist insbesondere $f \circ g : [a, b] \rightarrow \mathbb{R}$ definiert. Dann gilt:

$$\int_a^b f(g(x)) \cdot g'(x) \, dx = \int_{g(a)}^{g(b)} f(y) \, dy.$$

Beweis: Sei F eine Stammfunktion von f .

$$\begin{aligned} \int_{g(a)}^{g(b)} f(y) \, dy &= F(g(b)) - F(g(a)) \\ &= \int_a^b (F \circ g)'(x) \, dx \\ &= \int_a^b F'(g(x)) \cdot g'(x) \, dx. \end{aligned}$$

Dabei haben wir in den ersten beiden Schritten den Fundamentalsatz 4.10.21 und im letzten Schritt die Kettenregel verwendet. $\square$

Beispiel 4.10.28. (Anwendung der Substitutionsregel „von links nach rechts“)

1. Seien $0 < x_1, x_2$ und $\lambda > 0$. In der folgenden Rechnung setzen wir $f(y) = \frac{1}{y-1}$ und $g(x) = e^{\lambda x}$.

$$\begin{aligned} \int_{x_1}^{x_2} \frac{e^{\lambda x}}{e^{\lambda x} - 1} \, dx &= \frac{1}{\lambda} \int_{x_1}^{x_2} \underbrace{\frac{1}{e^{\lambda x} - 1}}_{f(g(x))} \cdot \underbrace{\lambda e^{\lambda x}}_{g'(x)} \, dx \\ &= \frac{1}{\lambda} \int_{e^{\lambda x_0}}^{e^{\lambda x_1}} \frac{1}{y - 1} \, dy \\ &= \frac{1}{\lambda} \ln(y - 1) \Big|_{e^{\lambda x_0}}^{e^{\lambda x_1}} \\ &= \frac{1}{\lambda} \ln(e^{\lambda x} - 1) \Big|_{x_1}^{x_2}. \end{aligned}$$

Also

$$\int \frac{e^{\lambda x}}{e^{\lambda x} - 1} \, dx = \frac{1}{\lambda} \ln(e^{\lambda x} - 1).$$

2. Wir berechnen nun eine Stammfunktion von $\tan x$ im Bereich $] \frac{-\pi}{2}, \frac{\pi}{2} [$. Dazu setzen wir $f(y) = \frac{1}{y}$ und $g(x) = \cos x$. Man beachte, dass in dem betrachteten Bereich $\cos x > 0$ gilt.

$$\begin{aligned} \int \tan x \, dx &= \int \frac{\sin x}{\cos x} \, dx \\ &= - \int \underbrace{\frac{1}{\cos x}}_{f(g(x))} \underbrace{(-\sin x)}_{g'(x)} \, dx \\ &= - \int \frac{1}{y} \, dy \Big|_{y=\cos x} \\ &= -\ln y|_{y=\cos x} + c \\ &= -\ln(\cos x) + c. \end{aligned}$$

Dabei ist die Notation in den beiden vorletzten Zeilen im Sinne von (4.34) in Bemerkung 4.10.22 zu verstehen.

Beispiel 4.10.29. (Anwendung der Substitution „von rechts nach links“)

1. Zunächst einmal schreiben wir einen häufig anzutreffenden Spezialfall der Substitutionsregel in einer etwas anderen Form auf, die insbesondere auch als praktische Merkhilfe dienen soll. Zur Berechnung von $\int_{y_0}^{y_1} f(y) dy$ substituieren wir die Variable y gemäß einer invertierbaren Transformation g :

$$y = g(x), \quad (4.36)$$

$$g^{-1}(y) = x. \quad (4.37)$$

Die Gleichung für die Ableitung $\frac{dy}{dx} = g'(x)$ schreiben wir *formal*

$$dy = g'(x) dx. \quad (4.38)$$

Desweiteren bemerken wir, welchen Integrationsgrenzen für x solche von y entsprechen:

$$y = y_i \Leftrightarrow x = g^{-1}(y_i) \quad \text{für } i = 1, 2. \quad (4.39)$$

Wir ersetzen nun *formal* in dem zu berechnenden Integral die Variable y durch $g(x)$, den Ausdruck dy durch $g'(x) dx$ und die Integrationsgrenzen y_i durch $g^{-1}(y_i)$ und erhalten so die Substitutionsregel für den Spezialfall einer invertierbaren Transformation g :

$$\int_{y_0}^{y_1} f(y) dy = \int_{g^{-1}(y_0)}^{g^{-1}(y_1)} f(g(x)) \cdot g'(x) dx. \quad (4.40)$$

Dies können wir als Regel zur Berechnung von Integralen ohne explizit gegebene Integrationsgrenzen schreiben:

$$\int f(y) dy = \int f(g(x)) \cdot g'(x) dx \Big|_{x=g^{-1}(y)}. \quad (4.41)$$

2. Wir berechnen erneut eine Stammfunktion zu $\ln x$ (vgl. Beispiel 4.10.26.2). Diesmal benutzen wir die uns bekannte Umkehrfunktion zu $\ln x$.

$$\begin{aligned} x &= \ln y, \\ y &= g(x) = e^x, \\ g'(x) &= e^x, \\ dy &= e^x dx. \end{aligned} \quad (4.42)$$

Wir substituieren also einfach den gesamten Integranden und integrieren partiell:

$$\begin{aligned} \int \ln y dy &= \int x \cdot e^x dx \Big|_{x=\ln y} \\ &= \left(x e^x - \int e^x dx \right) \Big|_{x=\ln y} \\ &= (x e^x - e^x) \Big|_{x=\ln y} \\ &= y \ln y - y. \end{aligned}$$

3. Im folgenden Beispiel möchten wir $\int \ln^2 y \, dy$ berechnen. In der Hoffnung, den komplizierten verketteten Ausdruck zu vereinfachen, wählen wir die Inverse der inneren Funktion als Transformation, also die gleich Substitution (4.43) wie im vorherigen Beispiel. Diese Identitäten verwenden wir in der folgenden Rechnung für die Substitutionen in (4.43). Von (4.43) auf (4.44) kommt man z.B. durch zweimalige partielle Integration, analog zu Beispiel 4.10.26.1.

$$\int \ln^2 y \, dy = \int x^2 e^x \, dx \Big|_{x=\ln y} \quad (4.43)$$

$$= x^2 e^x - 2x e^x + 2e^x \Big|_{x=\ln y} \quad (4.44)$$

$$= y \ln^2 y - 2y \ln y + 2y. \quad (4.45)$$

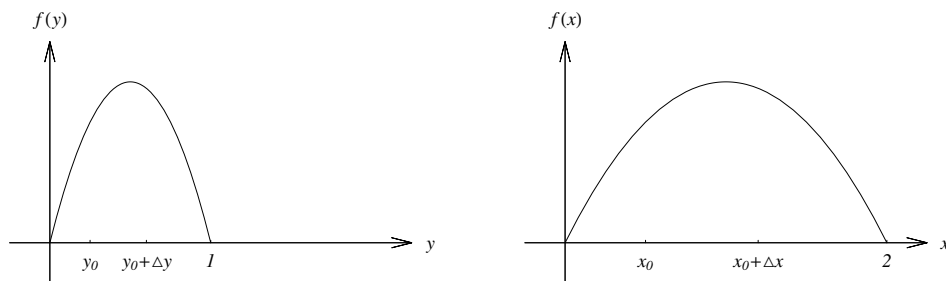


Abbildung 4.29: Streckung der Fläche bei Variablentransformation $y = \frac{1}{2}x$

Bemerkung 4.10.30. (Geometrische Bedeutung der Substitutionsregel)

Die formale Substitution $dy = g'(x)dx$ lässt sich auch geometrisch veranschaulichen. Dazu betrachten wir folgendes einfache Beispiel der Substitution

$$\begin{aligned} y &= g(x) = \frac{1}{2}x, \\ \Leftrightarrow x &= 2y, \\ dy &= \frac{1}{2}dx, \end{aligned}$$

welche wir wie folgt anwenden.

$$\int_0^1 f(y) \, dy = \int_0^2 f\left(\frac{1}{2}x\right) \cdot \frac{1}{2} \, dx$$

Durch die Substitution wird der Integrationsbereich gestreckt, und somit auch die Fläche, wie in Abbildung 4.29 illustriert. Damit die Integrale gleich sind, steht in dem neuen Integral das Reziproke dieses Streckfaktors. Allgemein gibt der Faktor $g'(x)$ an, wie stark der Integrationsbereich an der Stelle x (lokal) gestreckt wird, nämlich beim Übergang von der y -Koordinate auf

die x -Koordinate um den Faktor $\frac{1}{g'(x)}$. In der „mehrdimensionalen Integration“ wird das lokale Volumenverhältnis der Volumenelemente in den x - und den y -Koordinaten ebenfalls durch einen im Integral auftauchenden Faktor berücksichtigt, und zwar dem Absolutbetrag der Determinante der Jacobi-Matrix (erste Ableitung der Koordinatentransformation)

***Bemerkung 4.10.31. (für eine Funktion ohne elementar darstellbare Stammfunktion)**

Man kann, im Prinzip, beliebige durch elementare Funktionen (Polynome, e^x , $\sin x$ etc. und deren Umkehrfunktionen) dargestellte Funktionen systematisch differenzieren, d.h. durch (mechanisches) Anwenden der Differentiationsregeln erhält man für die erste Ableitung eine Darstellung durch elementare Funktionen.

Bei der analytischen Integration, d.h. dem Auffinden von Stammfunktionen, wie es hier gezeigt wurde, helfen oft, wenn überhaupt, nur scharfes Hinsehen und Probieren oder das Nachschlagen in Büchern mit Tabellen von Stammfunktionen oder ein entsprechendes mathematisches Computerprogramm zur analytischen Integration.

Es gibt allerdings auch integrierbare Funktionen, deren Stammfunktion sich *nicht elementar* darstellen lassen. Ein berühmtes Beispiel hierfür ist die Gaußsche Glockenkurvenfunktion (s. Abbildung 4.30)

$$f(x) = e^{-x^2}.$$

Die oben beschriebene Nicht-Darstellbarkeit der Stammfunktionen läßt sich in diesem Beispiel sogar mathematisch beweisen.

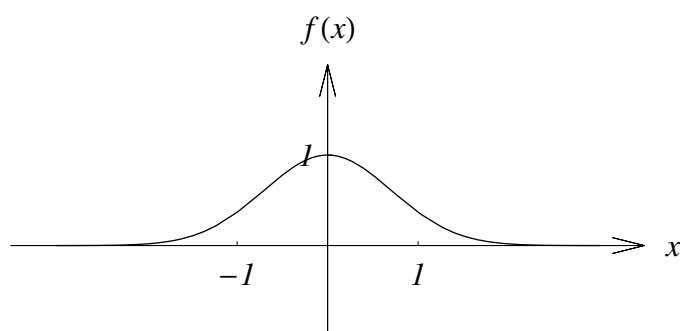


Abbildung 4.30: Graph der Funktion $f(x) = e^{-x^2}$

4.10.4 Uneigentliche Integrale

Bislang haben wir das Integral $\int_a^b f(x) dx$ nur für beschränkte Funktionen $f : [a, b] \rightarrow \mathbb{R}$ auf beschränkten Intervallen definiert. Was ist, wenn f oder der Integrationsbereich unbeschränkt sind? Wie kann man für solche Fälle die Definition des Integrals sinnvoll erweitern? Dazu wollen wir die zwei folgenden Beispiele betrachten.

Beispiel 4.10.32. (für unbeschränkte Integranden oder Integrationsbereiche)

1. (unbeschränkter Integrand)

$$\int_0^1 x^\alpha dx \quad \text{mit} \quad \alpha < 0. \quad (4.46)$$

Der Integrand ist auf $]0, 1]$ stetig, aber unbeschränkt und hat an der Stelle $x = 0$ eine Singularität.

2. (unbeschränkter Integrationsbereich)

$$\int_0^\infty e^{-x} dx. \quad (4.47)$$

Der Integrand ist beschränkt und stetig, der Integrationsbereich $[0, \infty[$ ist aber unbeschränkt.

Zunächst betrachten wir unbeschränkte Integranden mit genau einer Singularität auf einem beschränkten Integrationsbereich.

Definition 4.10.33. (uneigentliches Integral für singuläre Integranden)

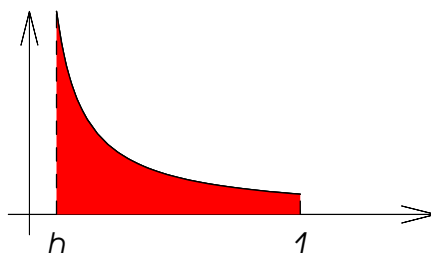
1. Sei $f : [a, b[\rightarrow \mathbb{R}$ und $\lim_{x \nearrow b} f(x) = \infty$. Wenn für jede Folge $(b_n)_{n \in \mathbb{N}}$ mit $a \leq b_n \leq b$ und $\lim_{n \rightarrow \infty} b_n = b$ der Grenzwert der Folge $\int_a^{b_n} f(x) dx$ existiert, dann definieren wir das **uneigentliche Integral** als

$$\int_a^b f(x) dx := \lim_{n \rightarrow \infty} \int_a^{b_n} f(x) dx. \quad (4.48)$$

2. Das uneigentliche Riemann-Integral ist für die Fälle $\lim_{x \nearrow b} f(x) = -\infty$, $\lim_{x \searrow a} f(x) = \pm\infty$ analog zu 1. definiert.
3. Für den noch allgemeineren Fall von *endlich vielen Singularitäten* von f definieren wir das uneigentliche Integral, indem wir das Intervall $[a, b]$ so zerlegen, dass f auf jedem Teilintervall höchstens an einem der Ränder eine Singularität hat. Ist f dann auf jedem Teilintervall integrierbar, so definieren wir $\int_a^b f(x) dx$ als Summe dieser Integrale.

Bemerkung 4.10.34.

1. In Definition 4.10.33.1 ist insbesondere vorausgesetzt, dass die betrachteten Integrale $\int_a^{b_n} f(x) dx$ existieren.
2. Desweiteren folgt aus den Voraussetzungen insbesondere (nach einem Standardargument), dass der betrachtete Grenzwert der Integrale unabhängig von der Folge $(b_n)_{n \in \mathbb{N}}$ ist. Damit ist (4.48) tatsächlich wohldefiniert.
3. Das Adjektiv *uneigentlich* wird oft auch weggelassen.

Abbildung 4.31: Das Integral $\int_h^1 \frac{1}{x} dx$ (y -Achse gestaucht.)

zu Beispiel 4.10.32.1: Im folgenden sei stets $h > 0$. Wir machen eine Fallunterscheidung für den Parameter α des Integranden f_α .

1. Fall: $\alpha = -1$.

$$\begin{aligned} \int_h^1 \frac{1}{x} dx &= \underbrace{\ln 1}_{=0} - \underbrace{\ln h}_{>0} \\ &= \ln \frac{1}{h} \\ \lim_{h \searrow 0} \ln \frac{1}{h} &= \infty \end{aligned}$$

Die Menge der Flächenmaße über $[h, 1]$ (mit $h > 0$) ist nach oben unbeschränkt, d.h. die Fläche wird beliebig groß bei entsprechender Wahl von h . (Vgl. Abbildung 4.31)

Also ist die Funktion nicht integrierbar.

2. Fall: $\alpha < -1$. Dann gilt $x^\alpha \geq \frac{1}{x}$ für $x \in]0, 1]$, also nach Fall 1 und der Monotonie des Integrals:

$$\lim_{h \searrow 0} \int_h^1 x^\alpha = \infty.$$

Also ist f_α auch in diesem Fall nicht integrierbar.

3. Fall: $-1 < \alpha < 0$.

$$\begin{aligned} \int_h^1 x^\alpha dx &= \left. \frac{1}{\alpha+1} x^{1+\alpha} \right|_h^1 \\ &= \frac{1}{1+\alpha} - \frac{1}{1+\alpha} h^{1+\alpha}. \end{aligned}$$

Wegen

$$\lim_{h \searrow 0} h^{1+\alpha} = 0$$

gilt also

$$\lim_{h \rightarrow 0} \int_h^1 x^\alpha dx = \frac{1}{1+\alpha} < \infty.$$

Folglich ist f_α integrierbar auf $[0, 1]$.

In diesem Beispiel haben wir also gesehen, dass die Funktion $f(x) = x^\alpha$ genau dann über $[0, 1]$ integrierbar ist, wenn $\alpha > -1$.

Definition 4.10.35. (uneigentliches Integral für unbeschränkte Intervalle)

1. Eine Funktion $f : [a, \infty[\rightarrow \mathbb{R}$ heißt **uneigentlich integrierbar auf $[a, \infty[$** , wenn für jede Folge b_n mit $b_n > 0$ und $\lim_{n \rightarrow \infty} b_n = \infty$, die Funktion $f|_{[a, b_n]} : [a, b_n] \rightarrow \mathbb{R}$ integrierbar ist und die Folge $\int_a^{b_n} f(x) dx$ konvergiert. In diesem Fall definieren wir

$$\int_a^\infty f(x) dx := \lim_{b \rightarrow \infty} \int_a^b f(x) dx.$$

2. Analog zu 1. definieren $\int_{-\infty}^a f(x) dx$.

3. Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ heißt **uneigentlich integrierbar auf $\mathbb{R}$** , wenn sie auf $] - \infty, 0]$ und auf $[0, \infty[$ uneigentlich integrierbar ist. In diesem Fall definieren wir

$$\int_{-\infty}^\infty f(x) dx := \int_{-\infty}^0 f(x) dx + \int_0^\infty f(x) dx.$$

zu Beispiel 4.10.32.2: Es gilt

$$\begin{aligned} \int_0^b e^{-x} dx &= -e^{-x} \Big|_0^b \\ &= -e^{-b} + e^{-0} \\ &= -e^{-b} + 1. \end{aligned}$$

Wegen

$$\lim_{b \rightarrow \infty} (-e^{-b} + 1) = 1$$

ist $f(x) = e^{-x}$ integrierbar auf $[0, \infty)$.

Bemerkung 4.10.36 (Rechenregeln für uneigentliche Integrale).

Partielle Integration, Substitutionsregel und der Fundamentalsatz (s. Sätze 4.10.24, 4.10.27 und 4.10.21) übertragen sich auf uneigentliche Integrale, vorausgesetzt dass die auftretenden Integrale existieren und die neuen Integrationsgrenzen und Randterme als entsprechende Grenzwerte wohldefiniert sind.

Beispiel 4.10.37 (Partielle Integration eines uneigentlichen Integrals). Wir berechnen das folgende uneigentlich Integral durch partielle Integration mit der Rollenverteilung $f(x) = x$ und

$g(x) = e^{-x}$, also $f'(x) = 1$ und $g(x) = -e^{-x}$.

$$\int_0^{\infty} x \cdot e^{-x} dx = -x \cdot e^{-x} \Big|_0^{\infty} + \int_0^{\infty} e^{-x} dx \quad (4.49)$$

$$\begin{aligned} &= \int_0^{\infty} e^{-x} dx \\ &= -e^{-x} \Big|_0^{\infty} \quad (4.50) \\ &= -0 + 1 = 1. \end{aligned}$$

Dabei verschwinden in (4.49) die beiden Randterme. Für $x = 0$ ist das klar, und an der oberen Intervallgrenze ist der Grenzwert $\lim_{x \rightarrow \infty} (-x \cdot e^{-x}) = 0$. Ebenso verschwindet wegen $\lim_{x \rightarrow \infty} (-e^{-x}) = 0$ in (4.50) der Randterm an der oberen Integrationsgrenze.

Problem des Riemann-Integrals in einigen Anwendungen:

$$\text{i.a. } \lim_{u \rightarrow \infty} \int f_u(x) dx \neq \int \lim_{u \rightarrow \infty} f_u(x) dx$$

Hinreichende Bedingung für „ $=$ “: f_u stetig differenzierbar $\forall u \in \mathbb{N}$ (in Anwendungen oft zu restriktiv)

Ausweg: Lebesgue-Integral (s. u.: Maß- und Integrationstheorie).

4.10.5 *Maß- und Integrationstheorie (Lebesgue-Integral)

Wir wollen eines “Maß” für Mengen M definieren. Streng genommen muss man noch die Arten von Mengen einschränken, für die die Definition eines Maßes im u.a. Sinn sinnvoll möglich ist. Diese technischen Aspekte sollen uns hier aber nicht weiter interessieren.

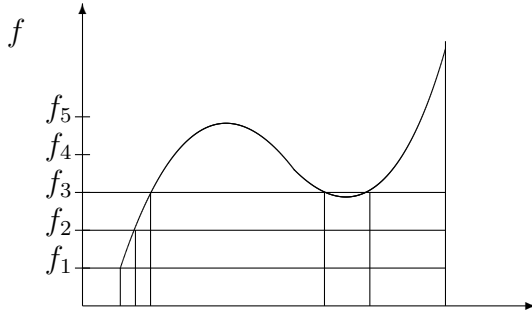
Definition 4.10.38. Sei $\mu : M \rightarrow \mathbb{R}$ eine Abbildung mit den Eigenschaften

- $\mu(\emptyset) = 0$ (leere Menge)
- $\mu(\bigcup_{k=1}^{\infty} M_k) = \sum_{k=1}^{\infty} \mu(M_k)$, falls $M_i, M_j \subset M, M_i \cap M_j = \emptyset \forall i \neq j$ (δ -Additivität)

Eine solche Abbildung μ nennt man ein Maß. Anwendungen finden sich z.B. in Wahrscheinlichkeitstheorie (Wahrscheinlichkeitsmaße). Ein Maß charakterisiert (misst) anschaulich gesprochen auf geeignete Art und Weise “die Größe” einer Menge. Damit wird es möglich, ein Integral über Funktionen auf messbaren Mengen zu definieren. Man beachte jedoch, dass man nicht für jede beliebige Menge ein Maß definieren kann.

Anschauliche Interpretation eines allgemeinen Integralbegriffs

Sei $f : M \rightarrow \mathbb{R}$ eine Abbildung, definiere $c_i := f_i$. Später dann $(f_{i+1} - f_i) \rightarrow 0$



Zur Einführung des allgemeinen Integralbegriffs wird also eine Intervalldiskretisierung der y-Achse statt x-Achse (wie beim Riemann-Integral) vorgenommen, Idee zur Definition des neuen Integralbegriffs (Lebesgue-Integral): „Sammeln“ der zu den jeweiligen f_i gehörenden x-Achsenabschnitte und Summation.

Definiere nun für stückweise konstante Funktionen

$$f(x) := \begin{cases} c_j & \forall x \in M_j, j = 1, \dots, m \\ 0 & \text{sonst} \end{cases}$$

ein Integral

$$\int_M f d\mu := \sum_{j=1}^m c_j \cdot \mu(M_j) \quad (4.51)$$

$$M = \bigcup_{j=1}^m M_j, \quad M_i \cap M_k = \emptyset, \quad \forall i \neq k$$

→ Verallgemeinerung des Integralbegriffs auf eine sehr große Klasse von Funktionen möglich !

Vorteile des Lebesgue-Integral sind:

- unter sehr allgemeinen Voraussetzungen gilt für Funktionenfolgen $\{f_n\}$

$$\lim_{n \rightarrow \infty} \int f_n d\mu = \int \lim_{n \rightarrow \infty} f_n d\mu \quad (4.52)$$

- Integration nicht nur über $\mathbb{R}$ möglich, sondern über sehr allgemeinen Mengen von Objekten (Ereignisse und Wahrscheinlichkeiten, Funktionenräume → wichtig in der Quantenmechanik, ...)

Das Riemann-Integral ist Spezialfall des allgemeineren Lebesgue-Integrals für das Maß

$$\mu([a, b]) := |b - a| \quad (4.53)$$

(siehe später: Integration im $\mathbb{R}^n$)

4.11 Numerische Differentiation und Integration

4.11.1 Numerische Berechnung von Ableitungen mit Hilfe von „Finiten Differenzen“- Approximationen (sog. Diskretisierungen)

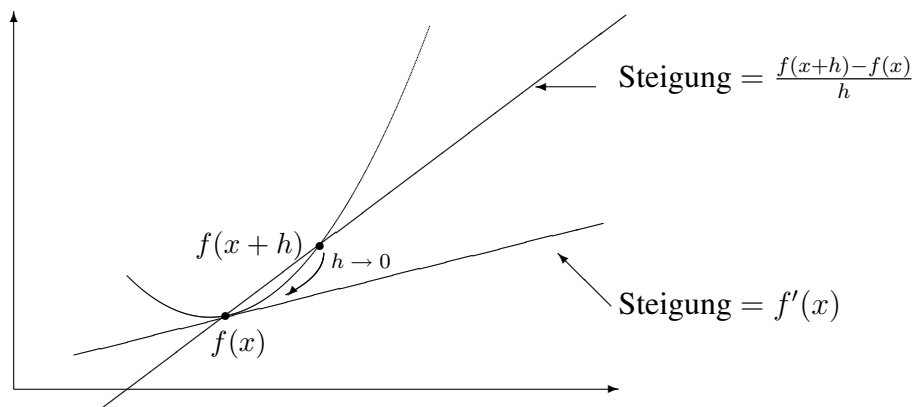
Es gilt

$$f'(x) = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} \approx \frac{f(x+h) - f(x)}{h}$$

für „kleine“ h („Vorwärtsdifferenzenquotient“).

Anschauliche Interpretation:

Nähere Tangente durch Sekanten an



Der Wert der Ableitung $f'(x)$ ist für kleines h gut durch die Sekantensteigung approximierbar.

Genauigkeit der Approximation

Taylor-Entwicklung von f um $f(x)$:

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

mit Restglied $R_{n+1}(x) = \frac{f^{(n+1)}(\xi)}{(n+1)!}(x+h-x)^{n+1}$ für ein geeignetes $\xi \in]x, x+h[$

Satz über Taylor-Reihe: Es gibt ein $\xi \in]x, x+h[$ mit

$$f(x+h) = f(x) + \sum_{i=1}^n \frac{f^{(i)}(x)}{i!} (x+h-x)^i + R_{n+1}(\xi), \quad h \text{ beliebig}$$

$$f(y) = f(x) + \sum_{i=1}^n \frac{f^{(i)}(x)}{i!} (y-x)^i$$

für $n = 1$ also:

$$f(x+h) = f(x) + f'(x) \cdot h + \underbrace{\frac{f''(\xi)}{2!} \cdot h^2}_{=:C} \Leftrightarrow \frac{f(x+h) - f(x)}{h} = f'(x) + C \cdot h$$

man schreibt dafür:

$$f'(x) = \frac{f(x+h) - f(x)}{h} + O(h)$$

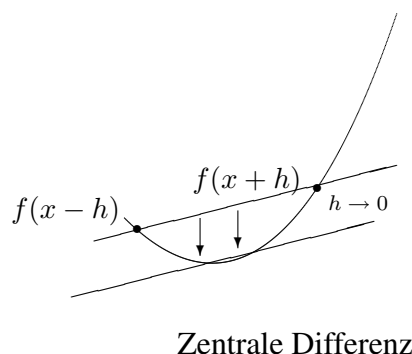
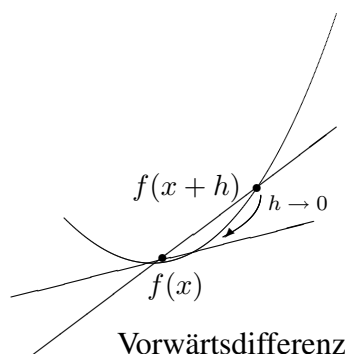
$O(h)$ meint $|C \cdot h|$. Man sagt auch die „Größenordnung“ h .

Die Vorwärtsdifferenz $\frac{f(x+h)-f(x)}{h}$ konvergiert für $h \rightarrow 0$ mit einer „Geschwindigkeit“ $O(h)$ gegen $f'(x)$. Die Approximation der Ableitung $f'(x)$ durch die Vorwärtsdifferenz ist also von der „Güte“ $O(h)$, bzw. der Fehler ist von der Größenordnung h , wobei h als „Schrittweite“ bezeichnet wird. Eine „Konvergenzordnung“ $O(h)$ heißt linear, $O(h^2)$ quadratisch.

Besser als Vorwärtsdifferenz: „Zentrale“ Differenz

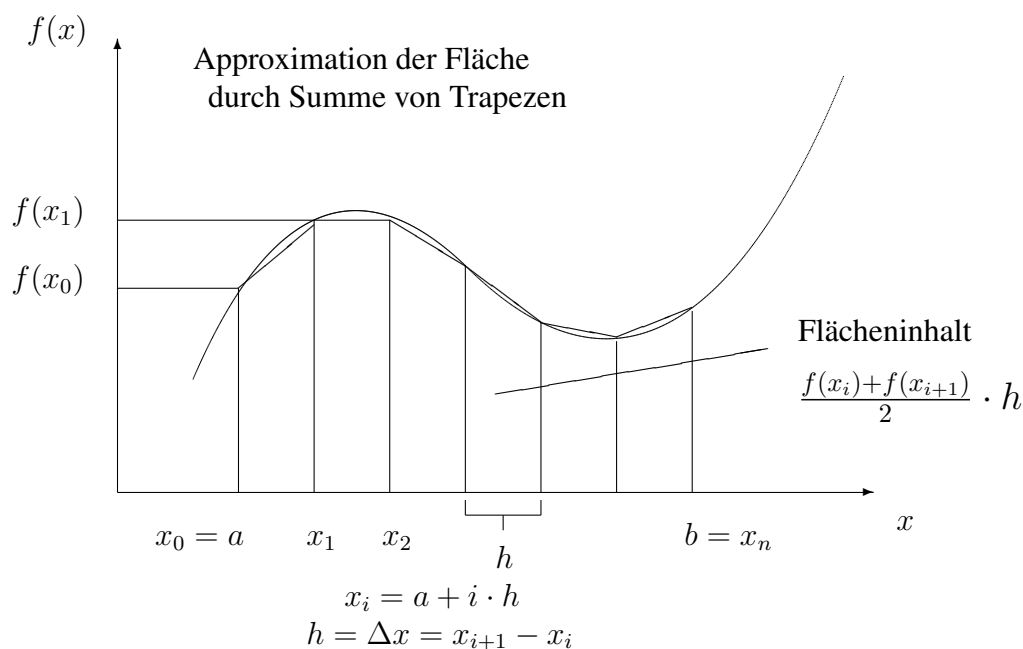
$$f'(x) \approx \frac{f(x+h) - f(x-h)}{2h} + O(h^2)$$

Beweis zur quadratischen Konvergenzordnung der zentralen Differenz: Übungsaufgabe



Alternativen

- AD (Automatische Differentiation)
 - „Zerlegung“ einer Funktion in Elementaroperationen, die der Computer zu ihrer Berechnung durchführt
 - Differentiation der Elementaroperation und geschicktes Anwenden der Kettenregel
- Symbolische Differentiation (Softwarepakete wie MAPLE oder MATHEMATICA)
 - man erhält einen expliziten analytischen Ausdruck für die Ableitung

4.11.2 Numerische Integration**Trapezregel**

$$\int_a^b f(x) dx \approx \sum_{i=0}^{n-1} \frac{f(x_i) + f(x_{i+1})}{2} \cdot h$$

Fehler der Trapezregel:

$$\left| \int_a^b f(x) dx - \sum_{i=0}^{n-1} \frac{f(x_i) + f(x_{i+1})}{2} \cdot h \right| \leq \frac{h^3}{12} \cdot \max_{x \in [a, b]} |f''(x)|$$

(sehr gute Approximation)

Kapitel 5

Komplexe Zahlen

Komplexe Zahlen werden das erstmal im 16. Jahrhundert beim Lösen von Gleichungen drittens Grades verwendet. Man führte hilfsweise Ausdrücke ein, die nicht als reelle Zahlen im herkömmlichen Sinne interpretiert werden konnten, und die man deshalb „imaginäre Zahlen“ nannte. Obwohl es zunächst viele Vorbehalte gegen diese seltsamen Objekte gab, überzeugten die verblüffende Eleganz und die vielen Erfolge beim Lösen praktischer Aufgaben im Laufe der Zeit alle Mathematiker von dem Sinn dieser Zahlen; an ihnen blieb jedoch noch lange etwas Mystisches haften; der Philosoph Gottfried Wilhelm Leibniz (1648-1716) schwärmte zum Beispiel: „Der göttliche Geist hat eine feine und wunderbare Ausflucht gefunden in jenem Wunder der Analysis, dem Monstrum der realen Welt, fast ein Amphibium zwischen Sein und Nicht-Sein, welches wir die imaginäre Einheit nennen.“ Heutzutage gehören die imaginären (bzw. komplexen Zahlen) zum Handwerkszeug nicht nur der Mathematiker und Physiker, sondern auch der Ingenieure und Chemiker, und natürlich auch der mathematischen Biologen. Mit ihrer Hilfe lassen sich viele Rechnungen leichter durchführen und wichtige Zusammenhänge besser verstehen.

5.1 Definition der Menge der komplexen Zahlen

Ausgehend von den reellen Zahlen nehmen wir die Zahl i (**die imaginäre Einheit**) mit der Eigenschaft

$$i^2 = -1, \quad (5.1)$$

hinzu und definieren die **Menge der komplexen Zahlen** durch

$$\mathbb{C} := \{x + iy \mid x, y \in \mathbb{R}\}.$$

Jede komplexe Zahl ist also durch ein Paar von reellen Zahlen gegeben. Für $z = x + iy$ bezeichnen wir

$$\begin{aligned} \operatorname{Re}(x + iy) &= x \quad \text{als } \mathbf{Realteil} \text{ von } z \text{ und} \\ \operatorname{Im}(x + iy) &= y \quad \text{als } \mathbf{Imaginärteil} \text{ von } z. \end{aligned}$$

Wir können uns $\mathbb{R}$ als Zahlengerade vorstellen und $\mathbb{C}$ als Ebene (s. Abbildung 5.1.) Komplexe Zahlen entsprechen dann Vektoren. Jeder Vektor in $\mathbb{C}$ kann durch seine Polarkoordinaten para-

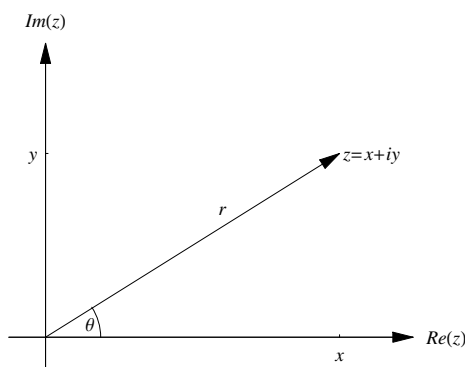


Abbildung 5.1: Die komplexe Zahlenebene

metrisiert werden.

$$\begin{aligned}
 z &= x + iy \\
 &= r(\cos \varphi + i \sin \varphi) \\
 &= re^{i\varphi}.
 \end{aligned}$$

In der letzten Gleichung haben wir die berühmte **Eulersche Identität** $\cos \varphi + i \sin \varphi = e^{i\varphi}$ verwendet, auf die wir an dieser Stelle aber nicht weiter eingehen (wer mag, kann ja einmal die Taylorreihe von $e^{i\varphi}$ mit der von $\sin \varphi$ und $\cos \varphi$ vergleichen).

Wir nennen r den **Absolutbetrag** (oder auch den **Betrag** oder den **Modul**) und φ das **Argument** von z . Der Betrag von z wird oft auch mit $|z|$ bezeichnet. Er ist die euklidische Länge des Vektors $(x, y) \in \mathbb{R}^2$.

Es gelten folgende Beziehungen:

$$x = r \cos \varphi, \quad (5.2)$$

$$y = r \sin \varphi, \quad (5.3)$$

$$r = |z| = \sqrt{x^2 + y^2}, \quad (5.4)$$

$$\varphi = \begin{cases} \arctan \frac{y}{x} & \text{für } x > 0, \\ +\frac{\pi}{2} & \text{für } x = 0, y > 0, \\ -\frac{\pi}{2} & \text{für } x = 0, y < 0, \\ \arctan \frac{y}{x} + \pi & \text{für } x \leq 0, y \geq 0, \\ \arctan \frac{y}{x} - \pi & \text{für } x < 0, y < 0. \end{cases} \quad (5.5)$$

5.2 Rechenregeln

Unter Verwendung von (5.1) können wir mit komplexen Zahlen so rechnen wie mit reellen. Zunächst betrachten wir *Addition*, *Subtraktion* und *Multiplikation*:

$$(x_1 + iy_1) + (x_2 + iy_2) = (x_1 + x_2) + i(y_1 + y_2) \quad (5.6)$$

$$(x_1 + iy_1) - (x_2 + iy_2) = (x_1 - x_2) + i(y_1 - y_2) \quad (5.7)$$

$$\begin{aligned} (x_1 + iy_1) \cdot (x_2 + iy_2) &= x_1x_2 + x_1 \cdot iy_2 + iy_1x_2 + iy_1 \cdot iy_2 \\ &= (x_1x_2 - y_1y_2) + i(x_1y_2 + y_1x_2) \end{aligned} \quad (5.8)$$

Addition und Subtraktion erfolgen also wie bei Vektoren und können entsprechend veranschaulicht werden (s. Abbildung 5.2.) Bei der Multiplikation haben wir (5.1) verwendet.

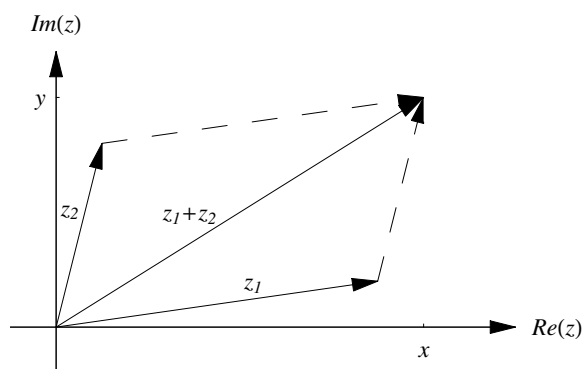


Abbildung 5.2: Addition von komplexen Zahlen

Mit Hilfe der Additionstheoreme für trigonometrische Funktionen können wir die Multiplikation von in Polarkoordinaten dargestellte komplexe Zahlen schreiben:

$$\begin{aligned} &(r_1(\cos \varphi_1 + i \sin \varphi_1)) \cdot (r_2(\cos \varphi_2 + i \sin \varphi_2)) \\ &= (r_1 \cdot r_2)((\cos \varphi_1 \cos \varphi_2 - \sin \varphi_1 \sin \varphi_2) + i(\cos \varphi_1 \sin \varphi_2 + \sin \varphi_1 \cos \varphi_2)) \\ &= (r_1 r_2)(\cos(\varphi_1 + \varphi_2) + i \sin(\varphi_1 + \varphi_2)). \end{aligned} \quad (5.9)$$

Die Absolutbeträge werden also multipliziert und die Argumente addiert modulo 2π , d.h. zur Summe der Argumente wird ein ganzzahliges Vielfaches von 2π addiert, sodass diese Summe im Intervall $(-\pi, \pi]$ liegt. S. Abbildung 5.3. Die **komplexe Konjugation** entspricht einer Spiegelung an der reellen Achse (s. Abbildung 5.4). Wir nennen $\bar{z}$ die zu z **konjugiert komplexe Zahl**.

$$x + iy = z \mapsto \bar{z} = x - iy, \quad (5.10)$$

$$r(\cos \varphi + i \sin \varphi) \mapsto r(\cos -\varphi + i \sin(-\varphi)) \quad (5.11)$$

$$= r(\cos \varphi - i \sin \varphi), \quad (5.12)$$

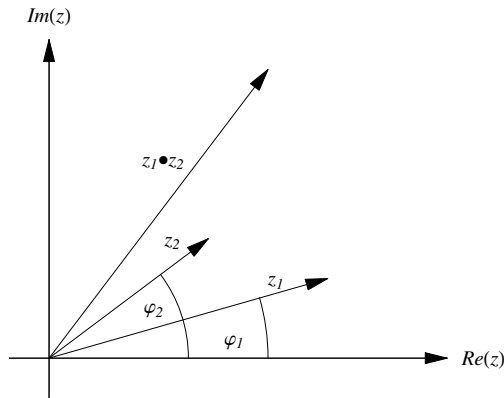


Abbildung 5.3: Multiplikation von komplexen Zahlen

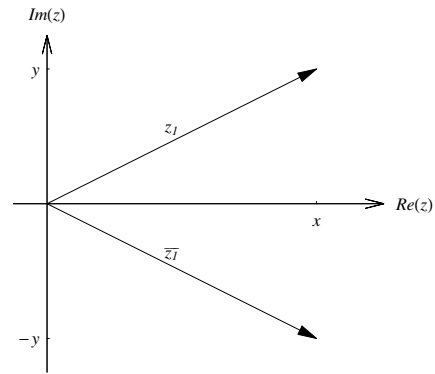


Abbildung 5.4: Konjugation einer komplexen Zahl

Satz 5.2.1. Seien $z = x + iy, z_1, z_2 \in \mathbb{C}$. Dann gilt:

$$\overline{\overline{z}} = z, \quad (5.13)$$

$$\overline{z_1 + z_2} = \overline{z_1} + \overline{z_2}, \quad (5.14)$$

$$\overline{z_1 \cdot z_2} = \overline{z_1} \cdot \overline{z_2}, \quad (5.15)$$

$$\operatorname{Re}(z) = \frac{z + \overline{z}}{2}, \quad (5.16)$$

$$\operatorname{Im}(z) = \frac{z - \overline{z}}{2i} \quad (5.17)$$

$$|z| = \sqrt{z\overline{z}} \quad (5.18)$$

$$|z| \geq 0 \quad (5.19)$$

$$|z| = 0 \Leftrightarrow z = 0 \quad (5.20)$$

$$|z_1 + z_2| \leq |z_1| + |z_2| \quad (\text{Dreiecksungleichung}) \quad (5.21)$$

Beweis: Die Aussagen (5.13) bis (5.18) folgen unmittelbar aus der Definition der komplexen Konjugation. Insbesondere ist die Zahl $x^2 + y^2 = z\overline{z}$ genau dann gleich 0, wenn $x = y = 0 \Leftrightarrow z = 0$, und ansonsten ist sie positiv. Somit ist die Wurzel (5.18) dieser Zahl eine wohldefinierte nicht-negative Zahl und genau dann gleich 0, wenn $z = 0$. Also gelten (5.19) und (5.20). Die Dreiecksungleichung (5.21) folgt aus der Dreiecksungleichung für $\mathbb{R}^2$. Man kann sie aber auch

leicht direkt zeigen:

$$\begin{aligned}
 |z_1 + z_2|^2 &= (z_1 + z_2) \cdot \overline{(z_1 + z_2)} = (z_1 + z_2)(\overline{z_1} + \overline{z_2}) \\
 &= z_1 \overline{z_1} + z_1 \overline{z_2} + z_2 \overline{z_1} + z_2 \overline{z_2} \\
 &= |z_1|^2 + z_1 \overline{z_2} + \overline{z_1} z_2 + |z_2|^2 \\
 &\leq |z_1|^2 + 2\operatorname{Re}(z_1 \overline{z_2}) + |z_2|^2 \\
 &\leq |z_1|^2 + 2|z_1||z_2| + |z_2|^2 \\
 &= (|z_1| + |z_2|)^2
 \end{aligned}$$

□

Bemerkung 5.2.2. Mit Hilfe von $|\cdot|$ läßt sich eine *Metrik* (Definition eines Abstandes zwischen zwei Punkten) auf $\mathbb{C}$ definieren.

$$d(z_1, z_2) := |z_1 - z_2|$$

Eine Metrik wird z.B. zur Definition von Konvergenz benötigt.

Wir berechnen das *multiplikativ Inverse* von $z \neq 0$, indem wir den Nenner reell machen, analog zum aus der Schule bekannten „Rational Machen“ von Nennern mit Wurzeltermen.

$$\frac{1}{z} = \frac{\overline{z}}{z\overline{z}} \quad (5.22)$$

$$= \frac{\overline{z}}{|z|^2}. \quad (5.23)$$

Mit der Darstellung $z = x + iy$ schreibt sich dies als

$$\frac{1}{z} = \frac{1}{x + iy} \quad (5.24)$$

$$= \frac{x - iy}{(x + iy)(x - iy)} \quad (5.25)$$

$$= \frac{x - iy}{x^2 + y^2} \quad (5.26)$$

$$= \frac{x}{x^2 + y^2} + i \frac{-y}{x^2 + y^2}. \quad (5.27)$$

In Polarkoordinaten erhalten wir

$$(r(\cos \varphi + i \sin \varphi))^{-1} = \frac{1}{r}(\cos \varphi - i \sin \varphi) \quad (5.28)$$

Geometrisch bedeutet die Abbildung $z \mapsto \frac{1}{z}$ die Inversion (Spiegelung) am Einheitskreis mit anschließender Spiegelung an der reellen Achse. (s. Abbildung 5.5.)

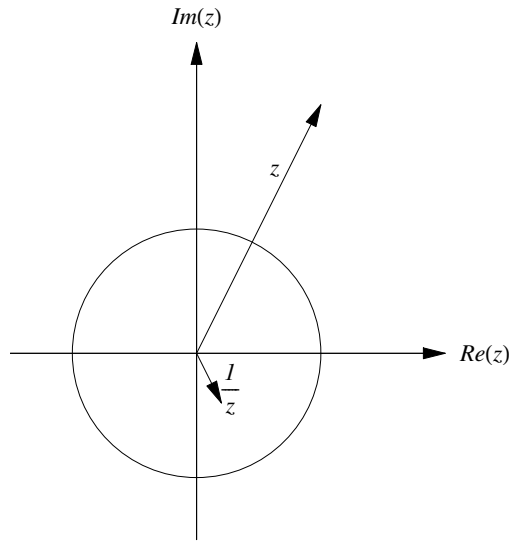


Abbildung 5.5: Inversion einer komplexen Zahl als Verknüpfung von Inversen am Einheitskreis und Spiegelung an der reellen Achse: $z \mapsto \frac{z}{|z|^2} = \frac{1}{\bar{z}} \mapsto \frac{1}{z}$.

Schliesslich können wir die *Division komplexer Zahlen* angeben (wobei wir auf eine Darstellung analog zu (5.6)-(5.8) verzichten):

$$\begin{aligned} \frac{z_1}{z_2} &= z_1 \cdot \frac{1}{z_2} \\ &= \frac{z_1 \bar{z}_2}{|z_2|^2}, \end{aligned} \quad (5.29)$$

oder in Polarkoordinaten:

$$\frac{r_1(\cos \varphi_1 + i \sin \varphi_1)}{r_2(\cos \varphi_2 + i \sin \varphi_2)} = \frac{r_1}{r_2}(\cos(\varphi_1 - \varphi_2) + i \sin(\varphi_1 - \varphi_2)) \quad (5.30)$$

d.h. die Beträge werden dividiert und die Argumente subtrahiert (modulo 2π).

5.3 Überblick über Zahlbereiche und deren Strukturen

Zum Abschluss dieses Kapitels geben wir in Tabelle 5.1 einen Überblick über die für uns wichtigsten Mengen von Zahlen und deren Strukturen. Der Übergang von einer Menge zur nächstgrößeren in unserer Liste wird dabei ganz pragmatisch motiviert. Wenn eine Menge bestimmte wünschenswerte Eigenschaften nicht besitzt (s. Spalte „Was geht nicht“), geht man zu einer größeren Menge mit dieser Eigenschaft über. Es können dabei allerdings auch Eigenschaften verlorengehen. Z.B. besitzen die komplexen Zahlen im Gegensatz zu den reellen keine Ordnung,

Menge	Struktur und Eigenschaften	„Was geht nicht“
ganze Zahlen $\mathbb{Z}$	1.) Ringstruktur, d.h. Verknüpfungen $+$, $-$ mit Axiomen. 2.) Totale Ordnung $<$, verträglich mit Ringstruktur.	Für $a \notin \{\pm 1\}$ hat die Gleichung $ax = 1$ keine Lösung, d.h. es gibt kein multiplikatives Inverses.
rationale Zahlen $\mathbb{Q}$	1.) $\mathbb{Q}$ ist ein Körper. 2.) Totale Ordnung $<$ verträglich mit Körperstruktur. 3.) Metrik: Abstand von $x_1, x_2 \in \mathbb{Q}$ ist $ x_1 - x_2 $.	$\mathbb{Q}$ ist bezüglich der Metrik nicht vollständig, d.h. $\mathbb{Q}$ hat Lücken. Bsp: Die Gleichung $x^2 - 2 = 0$ hat keine Lösung in $\mathbb{Q}$.
reelle Zahlen $\mathbb{R}$	1.) $\mathbb{R}$ ist ein Körper. 2.) Totale Ordnung verträglich mit Körperstruktur 3.) Metrik wie oben. 4.) $\mathbb{R}$ ist vollständig (s. Kapitel „Folgen“).	Die Gleichung $x^2 + 2 = 0$ hat keine reelle Lösung.
komplexe Zahlen $\mathbb{C}$	1.) $\mathbb{C}$ ist ein Körper. 2.) Metrik (s.o.), Vollständigkeit 3.) $\mathbb{C}$ ist algebraisch abgeschlossen, d.h. jedes nicht-konstante Polynom mit Koeffizienten aus $\mathbb{C}$ hat mindestens eine Nullstelle.	Keine Ordnung, die mit Körperstruktur verträglich ist.

Tabelle 5.1: Die Zahlbereiche $\mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R} \subset \mathbb{C}$

die mit der Körperstruktur verträglich ist. Die Erweiterungen der Mengen werden in Bemerkung 5.3.1 genauer erläutert.

Bemerkung 5.3.1.

1. In der Menge der ganzen Zahlen gibt es z.B. zu 2 kein multiplikativ inverses Element. Bsp.: Wenn man einen Kuchen gerecht auf zwei Leute verteilen möchte, dann erhalten beide mehr als nichts aber weniger als einen ganzen Kuchen, genauer gesagt jeder einen halben, also keinen ganzzahligen Anteil. Der Übergang von $\mathbb{Z}$ nach $\mathbb{Q}$ geschieht durch die Einführung von Brüchen, zusammen mit den bekannten Rechenregeln für diese.
2. Wie zu Beginn des Kapitels 4 erläutert, hat die Menge $\mathbb{Q}$ „Lücken“, wie durch das Beispiel der Lösung von $x^2 = 2$ erläutert wurde. Durch das „Stopfen“ dieser Lücken gelangt man von den rationalen zu den reellen Zahlen.
3. Gleichungen wie $x^2 = -1$ haben keine reelle Lösung. Durch die beschriebene Erweiterung der reellen zu den komplexen Zahlen werden insbesondere Lösungen solcher polynomiellen Gleichungen geschaffen. Im betrachteten Beispiel sind die beiden Lösungen i und $-i$.

Ganz wichtig für viele Bereiche der Mathematik ist der folgende Satz:

Satz 5.3.2. (Fundamentalsatz der Algebra)

Für jedes Polynom $p(x) = \sum_{k=0}^n a_k x^k$ mit Koeffizienten $a_k \in \mathbb{C}$, $a_n \neq 0$, gibt es n komplexe Zahlen $\bar{x}_1, \dots, \bar{x}_n$ (die „Nullstellen“ des Polynoms), so dass

$$p(x) = a_n(x - \bar{x}_1) \cdots (x - \bar{x}_n) \quad \forall x \in \mathbb{C}.$$

Kapitel 6

Skalarprodukte und Orthogonalität

6.1 Standard-Skalarprodukt in $\mathbb{R}^n$

Erinnerung: In Kapitel 2.4 im ersten Teil dieser Vorlesung wurde das **Standard-Skalarprodukt** im $\mathbb{R}^3$ eingeführt: Für $x, y \in \mathbb{R}^3$ ist

$$\langle x, y \rangle := x_1 y_1 + x_2 y_2 + x_3 y_3,$$

und x ist **orthogonal zu** y , wenn $\langle x, y \rangle = 0$. Die **euklidische Norm** oder auch **euklidische Länge** für Vektoren im $\mathbb{R}^3$ ist definiert durch

$$\begin{aligned}\|x\|_2 &= \sqrt{x_1^2 + x_2^2 + x_3^2} \\ &= \sqrt{\langle x, x \rangle}.\end{aligned}$$

Wir verallgemeinern nun diese Begriffe auf den Fall des Vektorraums $\mathbb{R}^n$.

Definition 6.1.1. (Standardskalarprodukt, Orthogonalität und euklidische Norm in $\mathbb{R}^n$)
Seien $x, y \in \mathbb{R}^n$. Wir definieren das **Standardskalarprodukt** durch

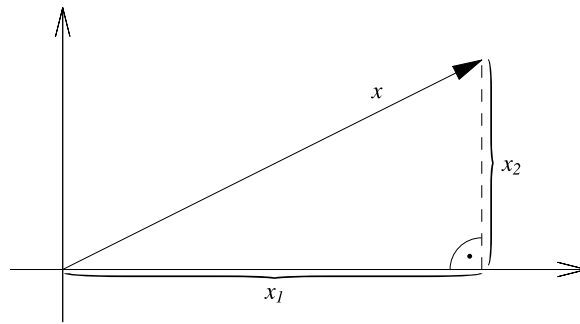
$$\langle x, y \rangle := x_1 y_1 + x_2 y_2 + \dots + x_n y_n.$$

Zwei Vektoren $x, y \in \mathbb{R}^n$ sind **orthogonal** zueinander, wenn $\langle x, y \rangle = 0$.

Die **euklidische Norm** oder auch **euklidische Länge** eines Vektors $x \in \mathbb{R}^n$ ist definiert als

$$\begin{aligned}\|x\|_2 &:= \sqrt{x_1^2 + x_2^2 + \dots + x_n^2} \\ &= \sqrt{\langle x, x \rangle}.\end{aligned}$$

Dabei ist die Definition der euklidischen Länge durch den Satz des Pythagoras motiviert. Für den Fall $n = 2$ vgl. Abbildung 6.1.

Abbildung 6.1: Vektor in $\mathbb{R}^2$

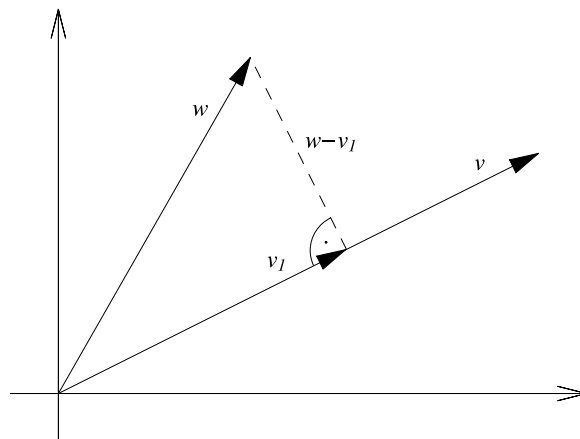
6.2 Orthogonale Projektion auf eine Gerade

Sei $V = \text{Spann}(v)$ ein eindimensionaler Untervektorraum des $\mathbb{R}^n$. Insbesondere gilt dann $v \neq 0$. Wir suchen zu einem Vektor $w \in \mathbb{R}^n$, der i.a. nicht in V liegt, die beste Approximation durch einen Vektor $v_1 \in V$. Diesen nennen wir auch das **Proximum** in V . Mathematisch präzisieren wir diese Aufgabe durch folgende

Problemstellung 6.2.1. (Minimierungsproblem: Proximum auf einer Geraden zu einem Punkt in $\mathbb{R}^n$)

Finde $v_1 \in V$, so dass $\|w - v_1\|_2$ minimal ist, also

$$\|w - v_1\|_2 = \min_{\tilde{v} \in V} \|w - \tilde{v}\|_2. \quad (6.1)$$

Abbildung 6.2: Das Proximum v_1 in $\text{Spann}(v)$ zu w

Durch Abbildung 6.2 motiviert, machen wir folgenden

Lösungsansatz: Wir wählen den Vektor v_1 so, dass $w - v_1$ orthogonal zu V ist. Wir ermitteln v_1 durch **orthogonale Projektion**. Zur Herleitung deren Berechnung formen wir die Bedingung,

dass der Vektor $w - v_1$ zu allen Vektoren aus $V = \{\lambda \cdot v \mid \lambda \in \mathbb{R}\}$ orthogonal ist, wie folgt um.

$$\begin{aligned}\langle w - v_1, \lambda v \rangle &= 0 \quad \forall \lambda \in \mathbb{R} \\ \Leftrightarrow \lambda \cdot \langle w - v_1, v \rangle &= 0 \quad \forall \lambda \in \mathbb{R} \\ \Leftrightarrow \langle w - v_1, v \rangle &= 0.\end{aligned}\tag{6.2}$$

Bemerkung 6.2.2. (Der Vorteil einer geometrischen Betrachtungsweise)

Gleichung (6.2) kann man lineares Gleichungssystem für die Koordinaten des Vektors v_1 auffassen. Wir gehen an dieser Stelle allerdings *nicht* zu der Koordinatendarstellung der Vektoren über. Dadurch erschweren wir uns nur den geometrischen (Durch-)Blick. Außerdem gelten folgende Rechnungen genauso für die orthogonale Projektion auf eine Gerade in einem beliebigen reellen Vektorraum mit Skalarprodukt (s. Definition 2.2.3.).

Da $v_1 \in V$, lässt es sich darstellen als

$$v_1 = \alpha \cdot v \quad \text{mit } \alpha \in \mathbb{R}.\tag{6.3}$$

Wir berechnen α , indem wir die Darstellung (6.3) in Gleichung (6.2) einsetzen.

$$\begin{aligned}0 &= \langle w - \alpha v, v \rangle \\ &= \langle w, v \rangle - \alpha \langle v, v \rangle \\ \Leftrightarrow \alpha &= \frac{\langle w, v \rangle}{\langle v, v \rangle}, \\ \text{also } v_1 &= \frac{\langle w, v \rangle}{\langle v, v \rangle} \cdot v.\end{aligned}\tag{6.4}$$

Wir empfehlen als Übung, zu überprüfen, dass $w - v_1$ mit v_1 aus (6.4) tatsächlich (6.2) erfüllt.

Satz 6.2.3 (Lösung des Minimierungsproblems).

Der in (6.4) definierte Vektor v_1 ist die eindeutige Lösung des Minimierungsproblems (6.1).

Beweis: Sei $v_2 \in V$ irgendein Vektor aus V . Wir können diesen schreiben als $v_2 = v_1 + v_3$ mit $v_3 = v_2 - v_1 \in V$. (S. Abbildung 6.3.) Der Vektor v_3 ist also gerade die Differenz von v_2 und v_1 . Wir berechnen das Quadrat der euklidischen Länge von $w - v_2$ und benutzen dabei die Rechenregeln für das Skalarprodukt sowie die Orthogonalität von $w - v_1$ zu v_3 .

$$\begin{aligned}\|w - v_2\|_2^2 &= \langle w - v_2, w - v_2 \rangle \\ &= \langle (w - v_1) - v_3, (w - v_1) - v_3 \rangle \\ &= \langle w - v_1, w - v_1 \rangle + \underbrace{\langle w - v_1, -v_3 \rangle}_{=0} \\ &\quad + \underbrace{\langle -v_3, w - v_1 \rangle}_{=0} + \langle -v_3, -v_3 \rangle \\ &= \underbrace{\|w - v_1\|_2^2}_{\text{konstant}} + \|v_3\|^2\end{aligned}\tag{6.5}$$

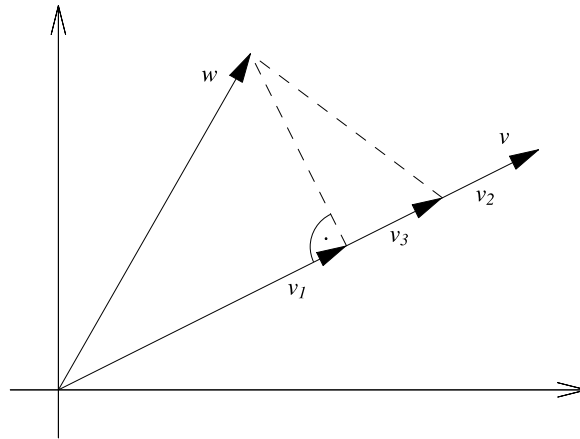


Abbildung 6.3: Zum Nachweis der Lösung des Minimierungsproblems

Das Quadrat der euklidischen Norm und somit die euklidische Norm von $(w - v_2)$ sind also genau dann minimal, wenn

$$\|v_3\|^2 = 0 \Leftrightarrow v_3 = 0 \Leftrightarrow v_2 = v_1.$$

Damit ist v_1 als die eindeutige Lösung von (6.1) nachgewiesen. □

Fazit: Wir erhalten das Proximum durch **orthogonale Projektion**:

$$w \mapsto \frac{\langle v, w \rangle}{\langle v, v \rangle} \cdot v =: P_V(w). \quad (6.6)$$

Falls v **normal** ist, d.h. $\|v\|_2 = \sqrt{\langle v, v \rangle} = 1$, dann vereinfacht sich (6.6) zu

$$P_V(w) = \langle v, w \rangle \cdot v. \quad (6.7)$$

Bemerkung 6.2.4 (Orthogonale Projektion als lineare Abbildung).

1. Die in (6.6) definierte Projektion ist eine lineare Abbildung

$$P_V : W \rightarrow V \subset W.$$

2. Für $w \in V$ gilt $P_V(w) = w$.
3. Der Koeffizient $\alpha = \langle v, w \rangle$ wird mit Hilfe des Skalarproduktes ausgerechnet.

Korollar 6.2.1 (Cauchy-Schwarz-Ungleichung).

Für alle $v, w \in \mathbb{R}^n$ gilt

$$|\langle w, v \rangle| \leq \|w\|_2 \cdot \|v\|_2, \quad (6.8)$$

und die Gleichheit in (6.8) gilt nur, falls w und v linear abhängig sind.

(Die Cauchy-Schwarz-Ungleichung gilt ganz allgemein für reelle Vektorräume mit Skalarprodukt (s. Definition 2.2.3.) Der Beweis dazu ist der gleiche.)

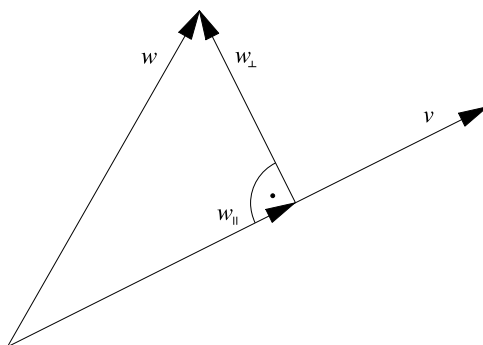


Abbildung 6.4: zum Beweis der Cauchy-Schwarz-Ungleichung: Zerlegung von w in eine zu v parallele Komponente $w_{\parallel}$ und ein zu v orthogonale $w_{\perp}$

Beweis: Falls $v = 0$, dann gilt offensichtlich die Gleichheit in (6.8).

Sei nun $v \neq 0$. Wir zerlegen w in eine zu v parallele und eine zu v orthogonale Komponente:

$$\begin{aligned} w &= w_{\parallel} + w_{\perp} \\ \text{mit } w_{\parallel} &:= \frac{\langle w, v \rangle}{\langle v, v \rangle} \cdot v, \\ w_{\perp} &:= w - \frac{\langle w, v \rangle}{\langle v, v \rangle} \cdot v. \end{aligned}$$

Diese beiden Komponenten sind orthogonal zueinander und somit gilt, analog zu (6.5),

$$\|w\|_2^2 = \|w_{\parallel}\|_2^2 + \|w_{\perp}\|_2^2.$$

Daraus erhalten wir die Abschätzungen

$$\begin{aligned} \|w\|_2^2 &\geq \left(\frac{\langle w, v \rangle}{\langle v, v \rangle} \right)^2 \cdot \|v\|_2^2 \\ &= \frac{(\langle w, v \rangle)^2}{\|v\|_2^4} \cdot \|v\|_2^2 \end{aligned} \tag{6.9}$$

$$\Leftrightarrow \|w\|_2 \cdot \|v\|_2 \geq |\langle w, v \rangle|. \tag{6.10}$$

wobei in (6.9) und (6.10) die Gleichheit nur gilt, wenn $w_{\perp} = 0$, d.h. wenn w und v linear abhängig sind. $\square$

Bemerkung 6.2.5 (Nicht-orientierter Winkel).

Aus der Cauchy-Schwarz-Ungleichung (6.8) folgt für zwei Vektoren $v, w \neq 0$:

$$-1 \leq \frac{\langle w, v \rangle}{\|v\|_2 \cdot \|w\|_2} \leq 1.$$

Dies ermöglicht uns, den **nicht-orientierten Winkel** $\angle(w, v)$ zwischen diesen beiden Vektoren zu definieren, und zwar durch

$$\cos(\angle(w, v)) := \frac{\langle w, v \rangle}{\|w\|_2 \cdot \|v\|_2}.$$

Auch diese Definition gilt wieder allgemein für reelle Vektorräume mit Skalarprodukt (s. Definition 6.4.1.) Diese Abstraktion wird sich als sehr nützlich erweisen, wenn wir in Bemerkung 9.1.52.1 in Kapitel 9.1.6 die Kovarianz als Skalarprodukt interpretieren.

6.3 Orthogonale Projektion auf einen Unterraum

Wir betrachten nun allgemein die orthogonale Projektion auf einen Untervektorraum des $\mathbb{R}^n$. Dazu sei ein **Orthogonalsystem** $(v_1, \dots, v_m)$ gegeben, d.h. $0 \neq v_i \in \mathbb{R}^n$ mit $\langle v_i, v_j \rangle = 0$ für $i \neq j$. Ein solches System ist insbesondere linear unabhängig

Beweis dazu: Sei $\alpha_1 v_1 + \dots + \alpha_m v_m = 0$ mit $\alpha_1, \dots, \alpha_m \in \mathbb{R}$. Dann gilt für jedes $1 \leq i \leq m$, dass $\alpha_i = 0$, wie wir durch die Bildung des Skalarproduktes beider Seiten der Vektorgleichung mit v_i sehen:

$$\begin{aligned} 0 &= \langle 0, v_i \rangle \\ &= \left\langle \sum_{l=1}^m \alpha_l v_l, v_i \right\rangle \\ &= \sum_{l=1}^m \alpha_l \underbrace{\langle v_l, v_i \rangle}_{=0 \text{ für } l \neq i} \\ &= \alpha_i \cdot \underbrace{\langle v_i, v_i \rangle}_{\neq 0 \text{ wegen } v_i \neq 0}. \end{aligned}$$

Das System $(v_1, \dots, v_m)$ spannt also einen m -dimensionalen Unterraum des $\mathbb{R}^n$ auf:

$$V = \text{Spann}(v_1, \dots, v_m) \subset \mathbb{R}^n.$$

Der folgende Satz ist eine Verallgemeinerung von Satz 6.2.3. In Abbildung 6.5 ist die orthogonale Projektion auf eine Ebene in $\mathbb{R}^3$ dargestellt.

Satz 6.3.1 (Orthogonale Projektion in $\mathbb{R}^n$).

Das Proximum zu $w \in \mathbb{R}^n$ in V ist durch orthogonale Projektion von w auf V gegeben, die man wie folgt berechnet:

$$P_V(w) = \sum_{i=1}^m \underbrace{\frac{\langle v_i, w \rangle}{\langle v_i, v_i \rangle}}_{\text{Koeffizient zu } v_i} \cdot v_i. \quad (6.11)$$

Falls die v_i normal sind, d.h. $\langle v_i, v_i \rangle = 1$, dann vereinfacht sich (6.11) zu

$$P_V(w) = \sum_{i=1}^m \langle v_i, w \rangle \cdot v_i. \quad (6.12)$$

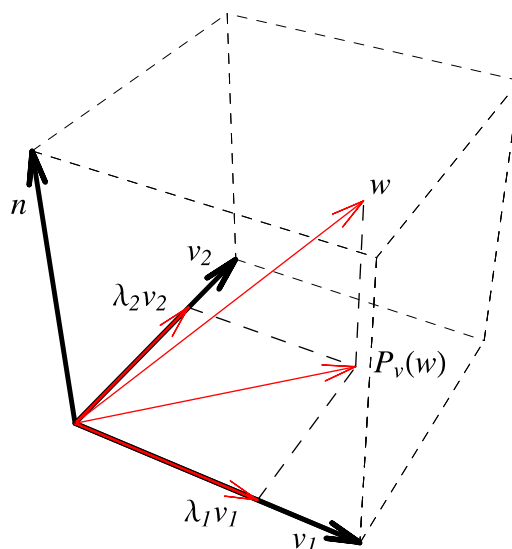


Abbildung 6.5: Orthogonale Projektion auf eine Ebene

Bemerkung 6.3.2. (Berechnung der Koeffizienten bzgl. einer Orthogonalbasis)

Die Koeffizienten von $P_V(w) \in V$ bezüglich der **Orthogonalbasis** $(v_1, \dots, v_m)$ von V werden einzeln durch Bildung von Skalarprodukten berechnet. Man muß kein lineares Gleichungssystem lösen wie z.B. sonst bei allgemeinen Basen (Koordinatensystemen). Dies macht den Gebrauch von Orthogonalbasen besonders attraktiv, insbesondere für effiziente numerische Berechnung bei praktischen Problemen. Siehe dazu auch Beispiel 6.6.3.

6.4 Skalarprodukte auf reellen Vektorräumen

Wir verallgemeinern noch einmal den Begriff des Skalarproduktes.

Definition 6.4.1 (Skalarprodukt auf einem reellen Vektorraum).

Sei W ein reeller Vektorraum. Ein **Skalarprodukt** auf W ist eine Abbildung

$$\langle \cdot, \cdot \rangle : W \times W \rightarrow \mathbb{R}$$

mit den folgenden Eigenschaften (Axiomen):

1. **(positive Definitheit)**

$$\begin{aligned} \forall w \in W \quad \langle w, w \rangle &\geq 0 \quad \text{und} \\ \langle w, w \rangle &= 0 \quad \Leftrightarrow w = 0. \end{aligned}$$

2. (Symmetrie)

$$\forall w_1, w_2 \in W \quad \langle w_1, w_2 \rangle = \langle w_2, w_1 \rangle.$$

3. (Linearität in beiden Argumenten)

$$\begin{aligned} \forall w_1, w_2, v \in W \quad \forall \alpha \in \mathbb{R} \quad \langle \alpha w_1 + w_2, v \rangle &= \alpha \langle w_1, v \rangle + \langle w_2, v \rangle \\ \langle v, \alpha w_1 + w_2 \rangle &= \alpha \langle v, w_1 \rangle + \langle v, w_2 \rangle. \end{aligned}$$

Das Skalarprodukt ist also eine **positiv definite, symmetrische Bilinearform**.

Beispiel 6.4.2. (für ein Skalarprodukt auf einem unendlich-dimensionalen Vektorraum)

Sei $W = \mathcal{C}^0([-\pi, \pi], \mathbb{R})$ der Raum der stetigen reellwertigen Funktionen auf dem Intervall $[-\pi, \pi]$. Zusammen mit der Addition von Funktionen und der Multiplikation von reellen Zahlen mit Funktionen bildet $\mathcal{C}^0([-\pi, \pi], \mathbb{R})$ einen unendlich-dimensionalen Vektorraum. Seine Elemente (Vektoren) sind Funktionen. Auf $\mathcal{C}^0([-\pi, \pi], \mathbb{R})$ definieren wir ein Skalarprodukt wie folgt. Seien $f, g \in \mathcal{C}^0([-\pi, \pi], \mathbb{R})$. Dann setzen wir

$$\langle f, g \rangle := \int_{-\pi}^{\pi} f(x) \cdot g(x) dx. \quad (6.13)$$

Wir bilden z.B. das Skalarprodukt der beiden Funktionen $f(x) = \sin x$ und $g(x) = 1$:

$$\begin{aligned} \langle f, g \rangle &= \int_{-\pi}^{\pi} (\sin x) \cdot 1 dx \\ &= 0. \end{aligned}$$

Also ist im Sinne des Skalarprodukts (6.13) die Sinusfunktion orthogonal zu jeder konstanten Funktion, was nichts anderes heißt, als dass Ihr Integral über dem Intervall $[-\pi, \pi]$ gleich 0 ist.

Definition 6.4.3 (Euklidische Norm).

Allgemein können wir mit Hilfe eines Skalarprodukts auf einem reellen Vektorraum W eine *Norm* (s. Definition 6.4.4) definieren. Für $w \in W$ setzen wir

$$\|w\|_2 := \sqrt{\langle w, w \rangle}.$$

Diese Norm heißt **die vom Skalarprodukt induzierte Norm** oder auch **euklidische Norm**.

Definition 6.4.4 (Norm auf einem reellen Vektorraum).

Sei W ein reeller Vektorraum. Eine Abbildung $\|\cdot\| : W \rightarrow \mathbb{R}$ heißt **Norm**, wenn folgende *Norm-Axiome* erfüllt sind:

1. (positive Definitheit)

$$\begin{aligned} \forall w \in W \quad \|w\| &\geq 0 \quad \text{und} \\ \|w\| &= 0 \Leftrightarrow w = 0. \end{aligned}$$

2. (Homogenität)

$$\forall w \in W \quad \forall \alpha \in \mathbb{R} \quad \|\alpha \cdot w\| = |\alpha| \cdot \|w\|.$$

3. (Dreiecksungleichung)

$$\forall w_1, w_2 \in W \quad \|w_1 + w_2\| \leq \|w_1\| + \|w_2\|.$$

Beispiel 6.4.5 (L_2 -Norm).

Die durch das Skalarprodukt (6.13) induzierte Norm auf $\mathcal{C}^0([-\pi, \pi], \mathbb{R})$ ist

$$\|f\|_2 := \left(\int_{-\pi}^{\pi} f(x) \cdot g(x) dx \right)^{\frac{1}{2}}. \quad (6.14)$$

Man nennt diese Norm die L_2 -Norm.

6.5 Fourier-Entwicklung

Wir betrachten wieder den Funktionenraum $\mathcal{C}^0([-\pi, \pi], \mathbb{R})$ und das Skalarprodukt (6.13) aus Beispiel 6.4.2. Zu diesem Raum definieren wir endlich-dimensionale Unterräume

$$V_n := \text{Spann} \left(\frac{1}{\sqrt{2\pi}}, \frac{1}{\sqrt{\pi}} \cos x, \dots, \frac{1}{\sqrt{\pi}} \cos(nx), \frac{1}{\sqrt{\pi}} \sin x, \dots, \frac{1}{\sqrt{\pi}} \sin(nx) \right)$$

Zwei Funktionen aus diesem aufspannenden System sind in Abbildung 6.6 dargestellt. Die

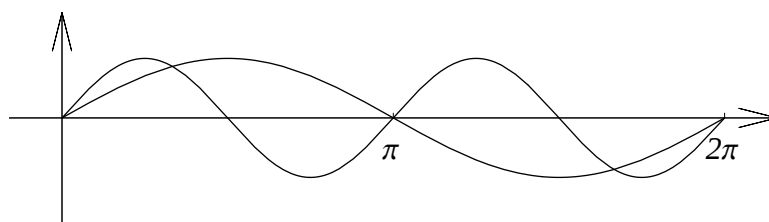
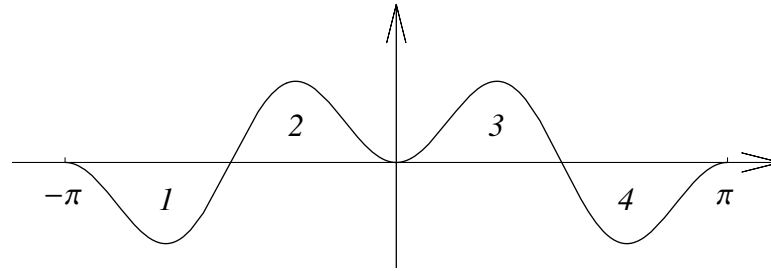


Abbildung 6.6: $\frac{1}{\sqrt{\pi}} \sin x$ und $\frac{1}{\sqrt{\pi}} \sin 2x$

Funktionen sind normiert und paarweise orthogonal, wie aus den unten stehenden Rechnungen hervorgeht, in denen $k \neq l$ gilt. Die hier zur Einübung der Integrationsregeln vorgeführte Berechnung der jeweiligen Stammfunktionen per Hand ist etwas mühsam. Es empfiehlt sich die Verwendung eines Computerprogramms mit symbolischer Rechnung oder das Nachschlagen der Stammfunktionen z.B. in [BSMM00]. Die hier betrachteten bestimmten Integrale lassen sich meist auch ohne Auffinden der Stammfunktion durch Ausnutzung von Punkt- und Achsensymmetrien der Integranden berechnen. Z.B. erkennt man in Abbildung 6.7 die Punktsymmetrie

Abbildung 6.7: $\sin x \cdot \sin(2x)$

der Funktion $f(x) = \sin x \sin(2x)$ bezüglich des Punktes $x = \frac{\pi}{2}$, die man auch schnell unter Verwendung der Symmetrien der Sinus-Funktion nachrechnen kann:

$$\begin{aligned}
 f\left(\frac{\pi}{2} - x\right) &= \sin\left(\frac{\pi}{2} - x\right) \sin\left(2\left(\frac{\pi}{2} - x\right)\right) \\
 &= \sin\left(\frac{\pi}{2} + x\right) \sin(\pi - 2x) \\
 &= \sin\left(\frac{\pi}{2} + x\right) \cdot (-1) \cdot \sin(\pi + 2x) \\
 &= -\sin\left(\frac{\pi}{2} + x\right) \sin\left(2\left(\frac{\pi}{2} + x\right)\right) \\
 &= -f\left(\frac{\pi}{2} + x\right).
 \end{aligned}$$

Aufgrund dieser Symmetrie addieren sich insbesondere die mit 3 und 4 markierten orientierten Flächeninhalte zu Null. Gleiches gilt für die Flächen 1 und 2. Also ist das Integral der Funktion f über dem Intervall $[-\pi, \pi]$ gleich Null.

Besonders elegant ist ein Beweis durch Integration der komplexwertigen Funktionen e^{ikx} und die Betrachtung von Real- und Imaginärteil, worauf wir hier aber nicht eingehen.

Nun kommen wir zu den angekündigten Rechnungen.

$$\begin{aligned}
 \left\langle \frac{1}{\sqrt{2\pi}}, \frac{1}{\sqrt{2\pi}} \right\rangle &= \frac{1}{2\pi} \int_{-\pi}^{\pi} 1 \, dx \\
 &= 1, \\
 \langle 1, \cos(kx) \rangle &= \left[\frac{\sin(kx)}{k} \right]_{-\pi}^{\pi} \\
 &= 0. \\
 \langle 1, \sin(kx) \rangle &= \left[-\frac{\cos(kx)}{k} \right]_{-\pi}^{\pi} \\
 &= 0.
 \end{aligned}$$

$$\begin{aligned}
\left\langle \frac{1}{\sqrt{\pi}} \cos(kx), \frac{1}{\sqrt{\pi}} \cos(kx) \right\rangle &= \frac{1}{\pi} \int_{-\pi}^{\pi} \sin^2(kx) dx \\
&= \frac{1}{\pi} \left[\frac{x}{2} + \frac{\sin(2kx)}{4k} \right]_{-\pi}^{\pi} \\
&= 1.
\end{aligned}$$

$$\begin{aligned}
\langle \cos(kx), \cos(lx) \rangle &= \int_{-\pi}^{\pi} \cos(kx) \cos(lx) dx \\
&= \left[\frac{\sin((k-l)x)}{2(k-l)} + \frac{\sin((k+l)x)}{2(k+l)} \right]_{-\pi}^{\pi} \\
&= 0.
\end{aligned}$$

$$\begin{aligned}
\left\langle \frac{1}{\sqrt{\pi}} \sin(kx), \frac{1}{\sqrt{\pi}} \sin(kx) \right\rangle &= \frac{1}{\pi} \int_{-\pi}^{\pi} \sin^2(kx) dx \\
&= \frac{1}{\pi} \left[\frac{x}{2} - \frac{\sin(2kx)}{4k} \right]_{-\pi}^{\pi} \\
&= 1.
\end{aligned}$$

$$\begin{aligned}
\langle \sin(kx), \sin(lx) \rangle &= \int_{-\pi}^{\pi} \sin(kx) \sin(lx) dx \\
&= \left[\frac{\sin((k-l)x)}{2(k-l)} - \frac{\sin((k+l)x)}{2(k+l)} \right]_{-\pi}^{\pi} \\
&= 0.
\end{aligned}$$

$$\begin{aligned}
\langle \cos(kx), \sin(kx) \rangle &= \int_{-\pi}^{\pi} \cos(kx) \sin(kx) dx \\
&= \left[\frac{-\cos^2(kx)}{2k} \right]_{-\pi}^{\pi} \\
&= 0.
\end{aligned}$$

$$\begin{aligned}
\langle \cos(kx), \sin(lx) \rangle &= \int_{-\pi}^{\pi} \cos(kx) \sin(lx) dx \\
&= \left[\frac{\cos((k-l)x)}{2(k-l)} - \frac{\cos((k+l)x)}{2(k+l)} \right]_{-\pi}^{\pi} \\
&= 0.
\end{aligned}$$

Wir können nun beliebige stetige Funktionen durch solche aus den Räumen V_n approximieren, analog zur Approximation durch orthogonale Projektion in (6.12).

$$P_{V_n}(f) = a_0 \cdot \frac{1}{\sqrt{2\pi}} + \sum_{k=1}^n a_k \frac{1}{\sqrt{\pi}} \cos(kx) + \sum_{k=1}^n b_k \cdot \frac{1}{\sqrt{\pi}} \sin(kx) \quad (6.15)$$

mit den **Fourier-Koeffizienten**

$$a_0 := \int_0^{\pi} f(x) \cdot \frac{1}{\sqrt{2\pi}} dx, \quad (6.16)$$

$$a_k := \int_0^{\pi} f(x) \cdot \frac{1}{\sqrt{\pi}} \cos(kx) dx \quad \text{für } k \geq 1, \quad (6.17)$$

$$b_k := \int_0^{2\pi} f(x) \cdot \frac{1}{\sqrt{\pi}} \sin(kx) dx \quad \text{für } k \geq 1. \quad (6.18)$$

Bemerkung 6.5.1 (Fourier-Koeffizienten).

In diesen Skript betrachten wir die orthonormalen Funktionen

$$\frac{1}{\sqrt{2\pi}}, \frac{1}{\sqrt{\pi}} \cos x, \dots, \frac{1}{\sqrt{\pi}} \cos(nx), \frac{1}{\sqrt{\pi}} \sin x, \dots, \frac{1}{\sqrt{\pi}} \sin(nx), \dots \quad (6.19)$$

und berechnen dazu die Koeffizienten gemäß (6.16)-(6.18). In der Literatur werden oft Systeme mit anders skalierten orthogonalen (nicht unbedingt normierten) Funktionen verwendet, z.B.

$$1, \cos x, \dots, \cos(nx), \sin x, \dots, \sin(nx), \dots$$

Dementsprechend erhält man andere Koeffizienten. Ebenso wird oft eine Fourier-Entwicklung auf anderen Intervallen betrachtet, z.B. auf $[0, 2\pi]$ oder auf $[0, 1]$, wobei für das letzte Intervall entsprechende orthogonale Funktionen $\dots, \cos(2\pi nx), \dots, \sin(2\pi nx), \dots$ verwendet werden müssen. Die Koeffizienten zu den hier genannten verschiedenen Systemen lassen sich leicht ineinander umrechnen, da man einen Vektor des einen Systems durch Skalierung eines entsprechenden Vektors aus dem anderen System erhält. (Das gilt natürlich i.a. nicht!) Wenn man z.B. aus einem Buch die Fourier-Koeffizienten einer Funktion übernimmt, sollte man darauf achten, zu welchem Funktionensystem sie gehören.

***Bemerkung 6.5.2** (Fourier-Reihe).

1. Im Grenzwert (für $n \rightarrow \infty$) erhält man die **Fourier-Reihe** oder **Fourier-Entwicklung** von f . Es gilt

$$\lim_{n \rightarrow \infty} \|f - f_n\|_2 = 0, \quad (6.20)$$

wobei wir die Notation $f_n := P_{V_n}(f)$ verwendet haben. Jedes $f \in C^0([-\pi, \pi], \mathbb{R})$ läßt sich im Sinne von (6.20) durch seine Fourier-Reihe darstellen, d.h. sich mit beliebiger Genauigkeit durch eine endliche Linearkombination von Vektoren des Systems (6.19) approximieren.

2. Wir bezeichnen das System in (6.19) daher auch als **vollständig**. Es ist also ein **vollständiges Orthonormalsystem**.
3. Die Fourier-Entwicklung existiert auch für beschränkte stückweise stetige Funktionen und es gilt (6.20). Gleichung (6.20) besagt die Konvergenz der Funktionenfolge bzgl. der in (6.14) definierten Norm. Auf andere Konvergenzgegriffe, z.B. punktweise Konvergenz (das hieße $f_n(x) \rightarrow f(x)$) gehen wir hier nicht ein.

Beispiel 6.5.3 (für eine Fourier-Reihe).

Wir berechnen die Fourier-Reihe der stückweise stetigen Funktion (s. Abbildung 6.8 und auch Abbildung 6.9)

$$f(x) = \begin{cases} -1 & \text{für } -\pi \leq x \leq 0, \\ 1 & \text{für } 0 < x < \pi. \end{cases} \quad (6.21)$$

Die Fourier-Koeffizienten sind

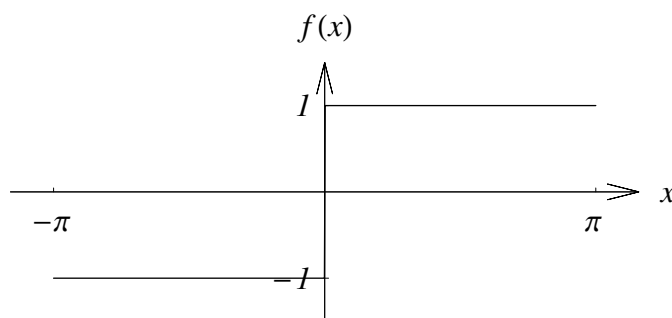


Abbildung 6.8: Stückweise konstante Funktion

$$\begin{aligned} a_0 &= \int_{-\pi}^{\pi} f(x) \cdot \frac{1}{\sqrt{2\pi}} dx \\ &= \frac{1}{\sqrt{\pi}} \left[\int_{-\pi}^0 (-1) dx + \int_0^{\pi} 1 dx \right] \\ &= 0. \end{aligned}$$

$$\begin{aligned}
\text{für } k \geq 1 : \quad a_k &= \frac{1}{\sqrt{\pi}} \int_{-\pi}^{\pi} f(x) \cos(kx) dx \\
&= \frac{1}{\sqrt{\pi}} \underbrace{\left[- \int_{-\pi}^0 \cos(kx) dx + \int_0^{\pi} \cos(kx) dx \right]}_{=0 \quad (\text{s.u.})} \quad (6.22)
\end{aligned}$$

$$\begin{aligned}
b_k &= \frac{1}{\sqrt{\pi}} \int_{-\pi}^{\pi} f(x) \sin(kx) dx \\
&= \frac{1}{\sqrt{\pi}} \left[- \int_{-\pi}^0 \sin(kx) dx + \int_0^{\pi} \sin(kx) dx \right] \quad (6.23)
\end{aligned}$$

$$= \frac{2}{\sqrt{\pi}} \int_0^{\pi} \sin(kx) dx \quad (6.24)$$

$$= \frac{2}{\sqrt{\pi}} \cdot \frac{1}{k} \int_0^{k\pi} \sin y dy \quad (6.25)$$

$$= \begin{cases} 0 & \text{für } k \text{ gerade,} \\ \frac{4}{\sqrt{\pi} \cdot k} & \text{für } k \text{ ungerade.} \end{cases} \quad (6.26)$$

Wir liefern nun einige Nebenrechnungen nach.

Der Term in eckigen Klammer in (6.22) ist gleich 0. Wir können nämlich den ersten Summanden durch die Substitution $x = -y \Leftrightarrow y = -x \Rightarrow dx = -dy$ wie folgt umformen.

$$\begin{aligned}
\int_{-\pi}^0 \cos(kx) dx &= \int_{\pi}^0 \cos(-ky) \cdot (-1) dy \\
&= - \int_{\pi}^0 \cos(ky) dy \\
&= \int_0^{\pi} \cos(ky) dy.
\end{aligned}$$

Im ersten Integralterm in (6.23) substituieren wir $x - y \Leftrightarrow y = -x \Rightarrow dx = -dy$:

$$\begin{aligned} - \int_{-\pi}^0 \sin(kx) dx &= - \int_{\pi}^0 \sin(-ky) dy \cdot (-1) dy \\ &= \int_0^{\pi} \sin(ky) dy \end{aligned}$$

und erhalten Zeile (6.24), in der wir vermöge $kx = y \Leftrightarrow x = \frac{1}{k}y \Rightarrow dx = \frac{1}{k}dy$ substituieren und so (6.25) erhalten. Von dort aus gelangen wir schließlich zu (6.26) durch die Überlegung, dass für natürliche Zahlen m Integrale der Form

$$\int_0^{2m\pi} \sin x dx = 0$$

verschwinden und so in (6.25) lediglich für ungerade $k = 2m + 1$ ein Integral

$$\int_{2m\pi}^{(2m+1)\pi} \sin x dx = 2$$

verbleibt. Insgesamt erhalten wir die Fourier-Reihe der Funktion f aus (6.21):

$$\frac{4}{\pi} \sum_{\substack{k=1, \\ k \text{ ungerade}}}^{\infty} \frac{1}{k} \sin(kx) = \frac{4}{\pi} \sum_{m=0}^{\infty} \frac{1}{(2m+1)} \sin((2m+1) \cdot x).$$

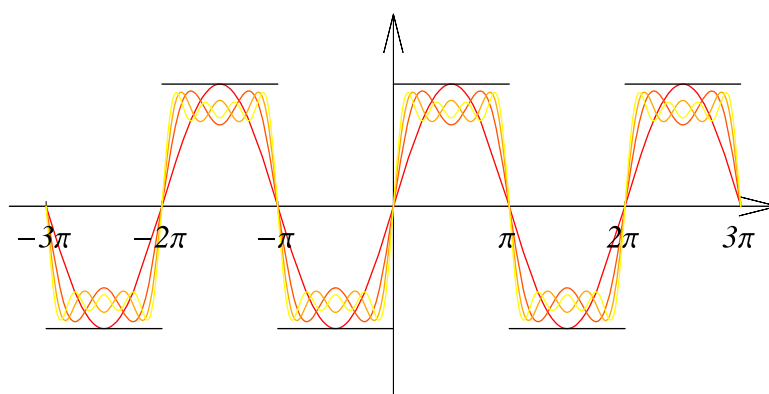


Abbildung 6.9: Die ersten Partialsummen f_n der Fourier-Reihe von f

Bemerkung 6.5.4 (Anwendung von Fourier-Reihen).

1. Eine praktische Anwendung der Fourier-Entwicklung ist ganz allgemein die **Analyse** von periodischen Signalen in ihre Frequenzanteile sowie die Erzeugung von periodischen Signalen aus Sinus-Schwingungen (**Synthese**), z.B. die Erzeugung einer elektronischen Sägezahn-Schwingung.
2. Auch theoretisch ist die Fourier-Entwicklung sehr wichtig, wie wir im übernächsten Abschnitt sehen werden.

6.6 *Orthonormalbasen und Selbstadjungierte Operatoren

In Kapitel 3.3.3 hatten wir schon auf die Vorteile der Diagonalisierbarkeit von Matrizen hingewiesen. Leider ist nicht jede Matrix diagonalisierbar, und man kann Matrizen im Allgemeinen auch nicht einfach ansehen, ob sie diagonalisierbar sind. Es gibt aber einige wichtige Spezialfälle, von denen wir zwei in diesem Abschnitt behandeln wollen, da sie für viele Bereiche der Physik und insbesondere für die theoretische Chemie sehr wichtig sind: Wir werden uns mit symmetrischen Matrizen beschäftigen, die man auch *selbstadjungiert* (bzw. im Komplexen *hermitesch*) nennt. Wir werden sehen, dass sie nicht nur diagonalisierbar sind, sondern dass die diagonalisierende Basistransformation sogar noch eine spezielle Struktur hat.

6.6.1 Orthonormalbasen und Orthogonale Matrizen

Die kanonischen Basisvektoren $e_1, \dots, e_n$ haben eine besonders schöne Eigenschaft, sie sind *orthogonal* zueinander (siehe Definition 6.1.1): Es gilt $\langle e_i, e_j \rangle = 0$ wenn $i \neq j$. Außerdem ist jeder Basisvektor e_i ein Einheitsvektor, d.h. er hat die Norm $\|e_i\| = 1$. Diese Eigenschaften der kanonischen Basis kann man auch bei anderen Basen feststellen, deren Basisvektoren wir uns als „gedrehte“ oder „gespiegelte“ Bilder der kanonischen Basisvektoren vorstellen können. Man nennt solche Basen „Orthonormalbasen“.

Definition 6.6.1 (Orthonormalbasis und Orthogonale Matrix).

Eine Basis $(v_1, \dots, v_n)$ eines Vektorraums mit Skalarprodukt (wie z.B. des $\mathbb{R}^n$) heißt **Orthonormalbasis**, wenn die Basisvektoren alle auf eins normiert sind und zueinander orthogonal sind, d.h. wenn gilt

$$\langle v_i, v_j \rangle = \delta_{ij} := \begin{cases} 1; & \text{wenn } i = j \\ 0; & \text{wenn } i \neq j. \end{cases} \quad (6.27)$$

Schreibt man im Falle des $\mathbb{R}^n$ die Basisvektoren als Spalten in eine Matrix $B := (v_1 | \dots | v_n)$, so ist diese Matrix **orthogonal**, d.h. es gilt $B^T B = \mathbb{I}_n$. Da B quadratisch ist, ist dies äquivalent zu $B^{-1} = B^T$.

Das sogenannte „Kronecker-Symbol“ δ_{ij} haben wir an dieser Stelle in (6.27) einfach einmal eingeführt, da es Ihnen in der Physik und Chemie möglicherweise wiederbegegnen könnte und die Notation manchmal sehr erleichtert. Man beachte, dass δ_{ij} einfach die Elemente der Einheitsmatrix darstellt, Einsen auf der Diagonalen ($i = j$), und sonst überall Nullen.

Koordinatentransformationen mit Orthonormalbasen sind besonders einfach: sind die Basisvektoren in der Matrix $B = (v_1 | \cdots | v_n)$, so erhält man die i -te Koordinate eines beliebigen Vektors y einfach durch Bilden des Skalarproduktes $\langle v_i, y \rangle$, und den gesamten Koordinatenvektor im neuen System durch Berechnen von $B^T y$. Es gilt die folgende Identität:

$$y = BB^T y = \left(v_1 \mid \cdots \mid v_n \right) \begin{pmatrix} v_1^T \\ \vdots \\ v_n^T \end{pmatrix} y = \sum_{i=1}^n v_i \langle v_i, y \rangle.$$

un man sieht, dass man y ganz einfach in seine „Komponenten“ $v_i \langle v_i, y \rangle$ zerlegen kann. Wir werden dies an zwei Beispielen verdeutlichen.

Beispiel 6.6.2. Die quadratische Matrix

$$B = (v_1 | v_2) = \left(\begin{array}{c|c} \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ -\frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{array} \right) = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ -1 & 1 \end{pmatrix}$$

ist orthogonal und ihre Spaltenvektoren v_1, v_2 formen eine Orthonormalbasis des $\mathbb{R}^2$. Wir prüfen dies leicht nach, indem wir die Skalarprodukte $\langle v_1, v_1 \rangle = 1$, $\langle v_1, v_2 \rangle = 0$ und $\langle v_2, v_2 \rangle = 1$ berechnen. Wie sehen nun aber die Koordinaten z.B. des Vektors $y = \begin{pmatrix} 10 \\ 1 \end{pmatrix}$ in dieser Basis aus?

Um den Koordinatenvektor $B^{-1}y$ in der neuen Basis zu erhalten, nutzen wir aus, dass $B^{-1} = B^T$, und berechnen einfach

$$B^T y = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & -1 \\ 1 & 1 \end{pmatrix} \begin{pmatrix} 10 \\ 1 \end{pmatrix} = \frac{1}{\sqrt{2}} \begin{pmatrix} 9 \\ 11 \end{pmatrix}.$$

Alternativ können wir diese Berechnung auch als

$$B^T y = \begin{pmatrix} v_1^T \\ v_2^T \end{pmatrix} y = \begin{pmatrix} v_1^T y \\ v_2^T y \end{pmatrix} = \begin{pmatrix} \langle v_1, y \rangle \\ \langle v_2, y \rangle \end{pmatrix} = \begin{pmatrix} \frac{(10-1)}{\sqrt{2}} \\ \frac{10+1}{\sqrt{2}} \end{pmatrix} = \begin{pmatrix} \frac{9}{\sqrt{2}} \\ \frac{11}{\sqrt{2}} \end{pmatrix}$$

interpretieren.

***Beispiel 6.6.3** (Haar-Basis, Datenkompression).

Die Vektoren

$$\begin{aligned}
 v_1 &= \frac{1}{\sqrt{8}} \begin{pmatrix} 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \end{pmatrix}^T \\
 v_2 &= \frac{1}{\sqrt{8}} \begin{pmatrix} 1 & 1 & 1 & 1 & -1 & -1 & -1 & -1 \end{pmatrix}^T \\
 v_3 &= \frac{1}{\sqrt{4}} \begin{pmatrix} 1 & 1 & -1 & -1 & 0 & 0 & 0 & 0 \end{pmatrix}^T \\
 v_4 &= \frac{1}{\sqrt{4}} \begin{pmatrix} 0 & 0 & 0 & 0 & 1 & 1 & -1 & -1 \end{pmatrix}^T \\
 v_5 &= \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & -1 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}^T \\
 v_6 &= \frac{1}{\sqrt{2}} \begin{pmatrix} 0 & 0 & 1 & -1 & 0 & 0 & 0 & 0 \end{pmatrix}^T \\
 v_7 &= \frac{1}{\sqrt{2}} \begin{pmatrix} 0 & 0 & 0 & 0 & 1 & -1 & 0 & 0 \end{pmatrix}^T \\
 v_8 &= \frac{1}{\sqrt{2}} \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 1 & -1 \end{pmatrix}^T
 \end{aligned}$$

bilden eine Orthonormalbasis des $\mathbb{R}^8$, was man leicht durch Prüfen der Normierung (z.B. $\langle v_2, v_2 \rangle = \frac{1}{8}(4 \cdot 1^2 + 4 \cdot (-1)^2) = 1$) und der Orthogonalität (z.B. $\langle v_2, v_8 \rangle = \frac{1}{\sqrt{8}\sqrt{2}}(6 \cdot 0 + (-1) \cdot 1 + (-1) \cdot (-1)) = 0$) bestätigen kann. In Abbildung 6.10 zeigen wir zur Veranschaulichung zwei der Basisvektoren figure Umgebung, ganze Breite, 2 Bilder Nebeneinander, 2 Captions, Diese Basis, die leicht auf höherdimensionale Räume verallgemeinert werden kann, wird auch

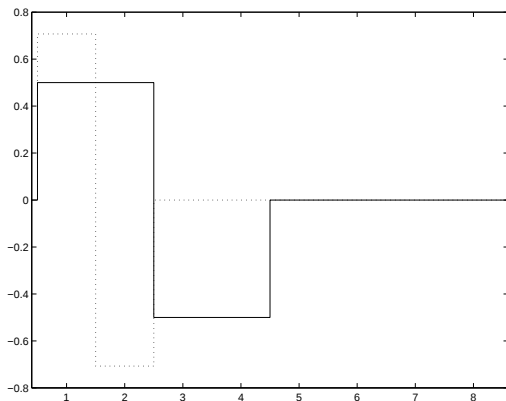


Abbildung 6.10: Die Basisvektoren v_3 (durchgezogene Linie) und v_5 (gepunktet) der Haar-Basis in Beispiel 6.6.3

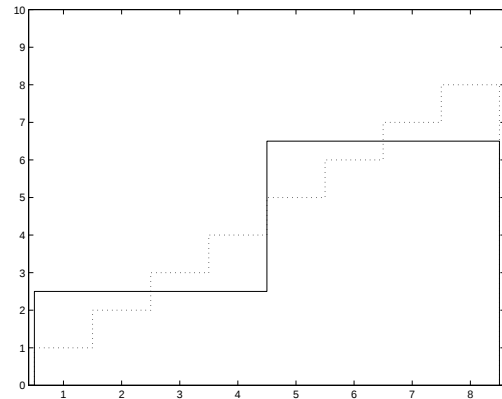


Abbildung 6.11: Die Approximation y' (durchgezogene Linie) durch die ersten zwei Komponenten, und der ursprüngliche Vektor y (gepunktet) aus Beispiel 6.6.3.

„Haar-Basis“ genannt (nach Alfred Haar, [Haa10]), und spielt besonders in der Datenkompression eine wichtige Rolle, wie wir gleich sehen werden. Zunächst berechnen wir, wie zuvor, die Koordinaten eines Vektors in der Basis $B = (v_1 | \dots | v_8)$; nehmen wir z.B. den Vektor

$$y = (1 \ 2 \ 3 \ 4 \ 5 \ 6 \ 7 \ 8)^T.$$

Wir bilden nun einfach nacheinander die Skalarprodukte $\langle v_i, y \rangle$ und erhalten die (gerundeten) Zahlenwerte

$$x := B^T y = (12.73 \ -5.66 \ -2 \ -2 \ -0.71 \ -0.71 \ -0.71 \ -0.71)^T.$$

Durch Bilden des Produkts Bx erhält man natürlich wieder den ursprünglichen Vektor y . Anstelle von y kann man sich also auch den Koordinatenvektor x merken. Beachten Sie, dass beide Vektoren aus 8 Zahlen bestehen.

Wie kann man die Haar-Basis nun zur Datenkompression nutzen? Man nutzt folgende Beobachtung: die hinteren Komponenten von x , die den „feineren“ Strukturen in y entsprechen, sind wesentlich kleiner als die ersten Komponenten – man könnte sie also, ohne einen großen Fehler zu machen, einfach weglassen und gleich Null setzen. Wenn wir uns also z.B. nur die ersten beiden Zahlen, x_1 und x_2 merken wollen, dann können wir den Vektor y statt durch den exakten Ausdruck

$$y = Bx = \sum_{i=1}^8 v_i x_i$$

auch durch die Approximation

$$y' = v_1 x_1 + v_2 x_2$$

ersetzen. Eine Veranschaulichung geben wir in Abbildung 6.11. Beachten Sie, dass man sich den Vektor y' mit Hilfe nur zweier Zahlen (x_1 und x_2) merken kann, während man sich für das exakte y alle 8 Komponenten merken muss.

Die Beobachtung, dass die „feineren“ Komponenten weniger Gewicht haben, also kleinere Koeffizienten in x , ist für sehr viele praktisch anfallende Daten erfüllt, zum Beispiel bei digitalisierten Bildern. Um solche Daten zu komprimieren, „dreht“ man sie einfach in eine Art Haar-Basis, und läßt dann die „feineren“ Komponenten weg. Man kann sich dann Bilder mit wesentlich weniger Zahlen merken, als sie Bildpunkte haben, unter leichtem Verlust der Bildauflösung. Man approximiert das ursprüngliche Bild also so, wie der Vektor y' mit nur zwei Zahlen den ursprünglichen Vektor y (der 8 Komponenten hat) approximiert. Praktische Rechnungen in höherdimensionalen Räumen (bei Bildern mit 600 mal 400 Bildpunkten arbeiten wir im $\mathbb{R}^{240000}$!) werden durch die Tatsache, dass die Basis orthonormal ist, überhaupt erst möglich.

6.6.2 Selbstadjungierte Operatoren und Symmetrische Matrizen

Eine quadratische reelle Matrix A heisst „symmetrisch“, wenn sie gleich ihrer Transponierten Matrix ist: $A = A^T$. Man kann diese Tatsache aber auch etwas abstrakter, mit Hilfe des Skalarproduktes, ausdrücken, und erhält dadurch neue interessante Einblicke. Lassen Sie sich nicht dadurch verwirren, dass wir statt „lineare Abbildung“ jetzt auch manchmal das gleichbedeutende Wort „Operator“ benutzen, um sie schonmal daran zu gewöhnen, dass Ihnen dieser Begriff besonders in der theoretischen Chemie noch häufiger begegnen wird.

Definition 6.6.4 (Selbstadjungierter Operator).

Ein Endomorphismus $f : V \rightarrow V$ in einem Vektorraum V mit Skalarprodukt (also z.B. der $\mathbb{R}^n$ mit dem Standard-Skalarprodukt) heisst „selbstadjungiert“ wenn für alle $v, w \in V$ gilt, dass

$$\langle f(v), w \rangle = \langle v, f(w) \rangle. \quad (6.28)$$

Der Begriff des selbstadjungierten Operators ist zwar allgemeiner als der einer symmetrischen Matrix, aber für unsere Zwecke sind sie fast identisch, denn:

Satz 6.6.5. Jede symmetrische Matrix $A = A^T \in \mathbb{R}^{n \times n}$ stellt einen selbstadjungierten Operator im $\mathbb{R}^n$ dar, und die darstellende Matrix A jedes selbstadjungierten Operators $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ im $\mathbb{R}^n$ (mit Standard-Skalarprodukt) ist symmetrisch.

Beweis: Seien $v, w \in \mathbb{R}^n$ beliebig. Dann ist Gleichung (6.28) für einen Operator f mit darstellender Matrix A äquivalent zu

$$v^T A^T w = (Av)^T w = \langle Av, w \rangle = \langle v, Aw \rangle = v^T Aw.$$

Damit ist bereits bewiesen, dass aus $A = A^T$ auch die Selbstadjungiertheit des dargestellten Operators folgt. Umgekehrt gilt, wenn wir $v = e_i$ und $w = e_j$ wählen, dass

$$a_{ij} = e_i^T A e_j = e_i^T A^T e_j = a_{ji},$$

d.h. die Matrix A muss symmetrisch sein, wenn nur die Selbstadjungiertheitsbedingung (6.28) erfüllt ist. $\square$

Für symmetrische Matrizen gilt nun der folgende und sehr erstaunliche Satz, der das wichtigste Ergebnis dieses Abschnittes ist.

Satz 6.6.6 (Diagonalisierbarkeit symmetrischer Matrizen).

Zu jeder symmetrischen Matrix $A \in \mathbb{R}^{n \times n}$ gibt es eine Orthonormalbasis $B = (v_1 | \cdots | v_n)$ des $\mathbb{R}^n$, die nur aus Eigenvektoren von A besteht, d.h. $D = B^T A B$ ist eine Diagonalmatrix. Außerdem sind alle Eigenwerte von A (also die Diagonalelemente von D) reell.

Für den sehr schönen Beweis dieses Satzes, den wir hier nicht vollständig wiedergeben, verweisen wir Interessierte auf Lehrbücher zur linearen Algebra, z.B. das Buch von Jähnich [Jäh98]. Um einen Geschmack der Beweistechnik zu bekommen, beweisen wir hier eine Teilaussage des Satzes.

Satz 6.6.7. (Orthogonalität von Eigenvektoren symmetrischer Matrizen zu verschiedenen Eigenwerten)

Sei A eine symmetrische Matrix, bzw. ein selbstadjungierter Operator, und seien v und w irgendzwei Eigenvektoren von A zu verschiedenen Eigenwerten. Dann sind v und w orthogonal zueinander.

Beweis: Sei

$$\begin{aligned} Av &= \lambda v, \\ Aw &= \mu w, \end{aligned}$$

mit $\lambda \neq \mu$. Es gilt

$$\begin{aligned}\lambda \langle v, w \rangle &= \langle \lambda v, w \rangle \\ &= \langle Av, w \rangle \\ &= \langle v, Aw \rangle \\ &= \langle v, \mu w \rangle \\ &= \mu \langle v, w \rangle.\end{aligned}$$

Daraus folgt

$$\begin{aligned}\underbrace{(\lambda - \mu)}_{\neq 0} \langle v, w \rangle &= 0 \\ \Rightarrow \langle v, w \rangle &= 0.\end{aligned}$$

□

Eine wichtige Anwendung des Satzes folgt in Abschnitt 6.6.4.

Beispiel 6.6.8. Wir betrachten als Beispiel eine zufällig erzeugte symmetrische Matrizen

$$A = \begin{pmatrix} 41 & 52 & 27 \\ 52 & 67 & 75 \\ 27 & 75 & 37 \end{pmatrix}$$

die wir in MATLAB bzw. SCILAB durch das Kommando `[B,D]=eig(A)` bzw. `[D,B]=bdiag(A)` diagonalisieren können, mit dem Ergebnis

$$B = \begin{pmatrix} 0.86102 & 0.24561 & 0.44531 \\ -0.17319 & -0.68168 & 0.71085 \\ -0.47816 & 0.68918 & 0.54441 \end{pmatrix} \quad \text{und} \quad D = \begin{pmatrix} 15.5462 & & \\ & -2.7561 & \\ & & 157.0149 \end{pmatrix}.$$

Man testet durch Eingabe von $B' * A * B$ leicht, dass tatsächlich wieder D herauskommt, und von $B' * B$, dass die Basis B tatsächlich orthonormal ist.

6.6.3 *Verallgemeinerung auf komplexe Matrizen

Für allgemeine Matrizen mit Elementen aus $\mathbb{C}$ heißt die Verallgemeinerung einer symmetrischen Matrix jetzt ganz einfach eine „selbst-adjungierte“ Matrix. Sie ist durch das Standard-Skalarprodukt im $\mathbb{C}^n$ definiert, das gegeben ist durch

$$\langle v, w \rangle = \sum_{i=1}^n \bar{v}_i w_i$$

wobei $\bar{z}$ wie zuvor in Kapitel 5 das komplex konjugierte einer komplexen Zahl z bezeichnet, und eine selbstadjungierte Matrix $A \in \mathbb{C}^{n \times n}$ muss dann einfach für alle $v, w \in \mathbb{C}^n$

$$\langle Av, w \rangle = \langle v, Aw \rangle$$

erfüllen. Man kann leicht zeigen, dass dies äquivalent ist zu $a_{ij} = \bar{a}_{ji}$. Wenn man im Komplexen arbeitet, benutzt man statt „selbst-adjungiert“ oft auch das Wort **hermitesch**. Man beachte, dass jede reelle symmetrische Matrix natürlich auch hermitesch ist, denn für reelle Einträge bleibt die komplexe Konjugation wirkungslos.

Die Eigenvektoren können nun aber sicher auch komplexe Einträge haben - wie können wir den Begriff der Orthonormalbasis bzw. den der orthogonalen Matrix verallgemeinern? Auch dies geschieht nun leicht mit Hilfe des Standard-Skalarproduktes im Komplexen, und eine Matrix $U = (v_1 | \dots | v_n) \in \mathbb{C}^{n \times n}$, die die Bedingung

$$\langle v_i, v_j \rangle = \delta_{ij} := \begin{cases} 1; & \text{wenn } i = j \\ 0; & \text{wenn } i \neq j \end{cases} \quad (6.29)$$

erfüllt, heisst nun **unitär**. Eine reelle orthogonale Matrix ist also auch unitär.

Für hermitesche Matrizen gilt nun der folgende Satz, der eine Verallgemeinerung von Satz 6.6.6 ist.

Satz 6.6.9 (Diagonalisierbarkeit hermitescher Matrizen).

Zu jeder hermiteschen Matrix $A \in \mathbb{C}^{n \times n}$ gibt es eine unitäre Matrix U , so dass $D = U^{-1}AU$ eine Diagonalmatrix ist. Außerdem sind alle Eigenwerte von A (also die Diagonalelemente von D) reell.

Wir beweisen hier wieder nur einen Teil des Satzes, nämlich dass die Eigenwerte reell sein müssen: Sei also λ ein Eigenwert von A und v der zugehörige Eigenvektor. Dann gilt:

$$\bar{\lambda} \langle v, v \rangle = \langle Av, v \rangle = \langle v, Av \rangle = \lambda \langle v, v \rangle$$

und wegen $\langle v, v \rangle \neq 0$ folgt $\bar{\lambda} = \lambda$, dass also λ reell sein muss. □

6.6.4 Der Laplace-Operator

Wir betrachten nun ein etwas abstrakteres Beispiel für einen selbstadjungierten Operator, das in der Physik von großer Bedeutung ist. Sei V der Raum der 2π -periodischen, beliebig oft differenzierbaren Funktionen. Für $f \in V$ sind auch alle Ableitungen $f^{(n)}$ von f Elemente von V : Aus $f(x + 2\pi) = f(x) \quad \forall x \in \mathbb{R}$ folgt nämlich durch n -maliges Ableiten und unter Verwendung der Kettenregel, dass $f^{(n)}(x + 2\pi) = f^{(n)}(x) \quad \forall x \in \mathbb{R}$.

Auf dem Vektorraum V ist die lineare Abbildung $\frac{-d^2}{dx^2}$, der **Laplace-Operator**, definiert:

$$\begin{aligned} \frac{-d^2}{dx^2} : V &\Rightarrow V \\ f &\mapsto -f''(x). \end{aligned}$$

Wir erwähnen, dass der Laplace-Operator manchmal auch als $\frac{d^2}{dx^2}$ definiert wird, also ohne das Minuszeichen. Dieser Differentialoperator ist natürlich allgemeiner auch auf zweimal-differenzierbare, nicht unbedingt 2π -periodische Funktionen anwendbar. Hier betrachten wir ihn jedoch nur als Operator auf dem speziellen Raum V . Die Funktionen

$$\frac{1}{\sqrt{\pi}}, \frac{1}{\sqrt{\pi}} \cos x, \frac{1}{\sqrt{\pi}} \cos(2x), \dots, \frac{1}{\sqrt{\pi}} \sin x, \frac{1}{\sqrt{\pi}} \sin(2x), \dots$$

sind Eigenvektoren des Laplace-Operators. Es gilt nämlich

$$\begin{aligned}
 -\frac{d^2}{dx^2} \left(\frac{1}{\sqrt{2\pi}} \right) &= 0 \\
 -\frac{d^2}{dx^2} \left(\frac{1}{\sqrt{\pi}} \cos x \right) &= \frac{1}{\sqrt{\pi}} \cos x \\
 &\vdots \\
 -\frac{d^2}{dx^2} \left(\frac{1}{\sqrt{\pi}} \cos(nx) \right) &= n^2 \cdot \frac{1}{\sqrt{\pi}} \cos(nx) \\
 &\vdots \\
 -\frac{d^2}{dx^2} \left(\frac{1}{\sqrt{\pi}} \sin(nx) \right) &= n^2 \cdot \frac{1}{\sqrt{\pi}} \sin(nx) \\
 &\vdots
 \end{aligned}$$

Der Laplace-Operator (definiert auf V) ist selbstadjungiert, d.h.

$$\left\langle \frac{d^2}{dx^2} f, g \right\rangle = \left\langle f, \frac{d^2}{dx^2} g \right\rangle \quad \forall f, g \in V.$$

Beweis dazu: Wir integrieren zweimal partiell. Die dabei auftretenden Randterme verschwinden wegen der 2π -Periodizität.

$$\begin{aligned}
 \left\langle \frac{d^2}{dx^2} f, g \right\rangle &= \int_{-\pi}^{\pi} (-f''(x)) \cdot g(x) dx \\
 &= \underbrace{[f'(x) \cdot g(x)]_{-\pi}^{\pi}}_{=0} + \int_{-\pi}^{\pi} f'(x) \cdot g'(x) dx \\
 &= \underbrace{[f(x) \cdot g'(x)]_{-\pi}^{\pi}}_{=0} - \int_{-\pi}^{\pi} f(x) \cdot g''(x) dx \\
 &= \int_{-\pi}^{\pi} f(x) \cdot (-g''(x)) dx \\
 &= \left\langle f, \frac{d^2}{dx^2} g \right\rangle.
 \end{aligned}$$

Ein selbstadjungierter Operator ist das Analogon zu einer **symmetrischen Matrix**, welche eine lineare Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^n$ darstellt, d.h. zu einer quadratischen Matrix A mit $A = A^T$. In Satz 6.6.6 hatten wir bereits gesehen, dass für solche Matrizen gilt, dass sie diagonalisierbar sind. Ganz analog gilt dies auch für jeden selbstadjungierten Operator, d.h. dass es eine Orthonormalbasis aus Eigenvektoren des Operators gibt. Für den Laplace-Operator ist die Fourier-Basis gerade diese Orthonormalbasis.

***Beispiel 6.6.10** (Die Wärmeleitungsgleichung).

Zur Modellierung der zeitlichen Entwicklung der Temperaturverteilung in einem dünnen kreisrunden Ring aus wärmeleitendem Material parametrisieren wir den Ring durch eine Winkelkoordinate x und beschreiben die Temperaturverteilung durch eine Funktion in x und der Zeitvariablen t , also

$$\begin{aligned} f : \mathbb{R}^{\geq 0} \times \mathbb{R} &\rightarrow \mathbb{R}, \\ (t, x) &\mapsto f(t, x). \end{aligned}$$

Also $f(t, x)$ ist die Temperatur zur Zeit t an der Stelle x . Für jedes t ist die durch $x \mapsto f(t, x)$ gegebene Funktion 2π -periodisch und beschreibt die Temperaturverteilung zur Zeit t . Für festes x beschreibt die Funktion $t \mapsto f(t, x)$ den zeitlichen Temperaturverlauf der an der Stelle x .

Zum Zeitpunkt $t = 0$ sei die Temperatur vorgegeben durch $f_0 \in V$. Wir stellen also die **Anfangsbedingung**

$$\forall x \in \mathbb{R} \quad f(0, x) = f_0(x). \quad (6.30)$$

Physikalisch ist die Temperatur nach unten beschränkt. Darauf gehen wir hier nicht weiter ein. Die zeitliche Entwicklung der Temperaturverteilung wird durch die Wärmeleitungsgleichung modelliert:

$$\forall (t, x) \in \mathbb{R}^{\geq 0} \times \mathbb{R} \quad \frac{\partial}{\partial t} f(t, x) = c \cdot \frac{\partial^2}{\partial t^2} f(t, x), \quad (6.31)$$

wobei die Konstante $c > 0$ die Wärmeleitfähigkeit des Materials beschreibt. Gleichung (6.31) ist eine **partielle Differentialgleichung**. Das **Anfangswertproblem**, gegeben durch (6.31), die Anfangsbedingung (6.30) und die Forderung der Differenzierbarkeit und Periodizität von f beschreibt die Umverteilung der Wärme durch Diffusion. Dabei bleibt die gesamte Wärmeenergie erhalten.

Wir bemerken, dass das betrachtete Problem stets eine Eindeutige Lösung hat. Auf die Existenz und Eindeutigkeit der Lösung gehen wir hier aber nicht näher ein.

Zur Illustration betrachten wir nun die jeweiligen Lösungen zu zwei verschiedenen Anfangsbedingungen, die jeweils Eigenwerte des Operators $c \cdot \frac{\partial^2}{\partial t^2}$ sind.

1. (konstante Anfangsverteilung)

Zur Anfangsbedingung

$$f_0(x) = 1$$

ist die Lösung des Anfangswertproblems

$$f(t, x) = 1,$$

da offensichtlich f die geforderten Differenzierbarkeits- und Periodizitätsbedingungen erfüllt und

$$\begin{aligned} f(0, x) &= f_0(x) \\ \frac{\partial}{\partial t} f(t, x) &= 0 \\ &= c \cdot \frac{\partial^2}{\partial t^2} f(t, x). \end{aligned}$$

Die konstante Temperaturverteilung ändert sich also nicht mit der Zeit. Das System befindet sich im (makroskopischen) Gleichgewicht.

2. (nicht-konstante Anfangsverteilung)

Die Lösung zur Anfangsbedingung

$$f_0(x) = \sin(nx)$$

ist

$$f(t, x) = e^{-cn^2 t} \sin(nx)$$

wie wir leicht überprüfen: Die Funktion f erfüllt die geforderten Differenzierbarkeits- und Periodizitätsbedingungen und außerdem die Anfangsbedingung, da

$$e^{-cn^2 \cdot 0} = 1,$$

und Gleichung (6.31):

$$\begin{aligned} \frac{\partial}{\partial t} f(t, x) &= -cn^2 \cdot f(t, x) \\ &= c \cdot \frac{\partial^2}{\partial x^2} f(t, x). \end{aligned}$$

Wir sehen, dass sich die Temperaturunterschiede mit der Zeit ausgleichen, und zwar exponentiell schnell mit der Rate cn^2 , welche bis auf ein Vorzeichen dem zum Eigenvektor f_0 des Differentialoperators $c \cdot \frac{\partial^2}{\partial x^2}$ gehörigen Eigenwert gleicht. Je größer n ist, also je stärker die Temperaturverteilung zu $t = 0$ oszilliert, desto größer ist diese Rate.

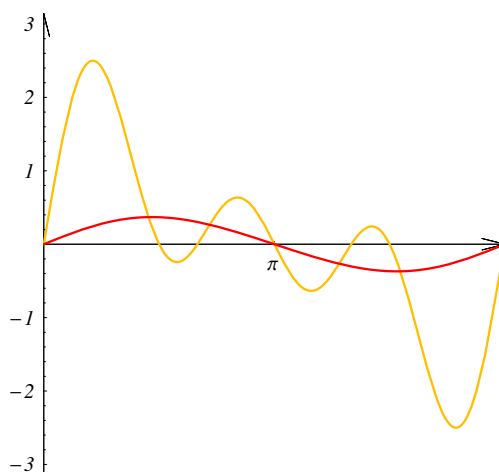


Abbildung 6.12: Zeitliche Entwicklung einer nicht-konstanten Temperaturverteilung

In beiden Fällen haben wir als Anfangsbedingung einen Eigenvektor (eine Eigenfunktion) des Differentialoperators $c \cdot \frac{\partial^2}{\partial x^2}$ betrachtet. Zu den Eigenvektoren lässt sich die Lösung recht einfach

darstellen. Wir erinnern uns an eine ähnliche Situation, und zwar bei Modell II zur Kaninchenpopulation im ersten Semester. Bei diesem ist die zeitliche Entwicklung eines Zustandes ebenfalls durch einen linearen Operator gegeben. Analog dazu können wir auch hier allgemeine Anfangszustände mit Hilfe von Eigenvektoren des linearen Operators darstellen (Analyse), nämlich durch ihre jeweilige Fourier-Reihe, dann für jede einzelne Fourier-Komponente das Problem lösen, d.h. die zeitliche Entwicklung berechnen, und diese schließlich wieder zusammensetzen (Synthese). Zur Illustration sind in Abbildung 6.12 die Anfangstemperaturverteilung $f(x, 0)$ und die Temperaturverteilung $f(x, 1)$ zur Zeit $t = 1$ abgebildet, mit $f(x, t) = \sum_{n=1}^3 e^{-n^2 t} \sin(nt)$. Der Koeffizient c in der Wärmeleitungsgleichung ist hier gleich 1.

Bezug zur Quantentheorie

In der Quantenmechanik (in der theoretischen Chemie) wird der Zustand eines Systems (z.B. Wasserstoff-Atom) durch eine komplexwertige Funktion beschrieben (**Wellenfunktion**). Auf Räumen solchen Funktionen werden **hermitesche Operatoren** (s. Abschnitt 6.6.3) betrachtet, die ein Analogon zu den selbstadjungierten Abbildungen auf reellen Vektorräumen darstellen. Zu diesen speziellen Operatoren (Hamilton-Operatoren, Drehimpuls-Operator etc.) werden Eigenvektoren (diese entsprechen den Orbitalen) berechnet. Die entsprechenden Eigenwerte werden **Quantenzahlen** genannt.

Kapitel 7

Analysis im $\mathbb{R}^n$

Wir haben bereits im Detail untersucht, wie man Ableitungen und Integrale von skalaren Funktionen $f : \mathbb{R} \rightarrow \mathbb{R}$ berechnet; dies sind allerwichtigste Basistechniken, die Ihnen immer wieder Hilfe leisten werden. In vielen Anwendungen tauchen jedoch Funktionen

$$f : \mathbb{R}^n \longrightarrow \mathbb{R}^m.$$

mit mehreren Argumenten und/oder mehreren Werten auf, und häufig möchte man auch hier mit Ableitungen oder Integralen arbeiten. Beispiele solcher Fragestellungen sind z.B.

- Wie bestimmt man die Länge einer Spirale?
- Wie berechnet man die Sonneneinstrahlung auf ein Blatt?
- Wie beschreibt man eine Flüssigkeitsströmung?
- Wann ist ein Gleichgewicht stabil?
- Welche Energie enthält ein elektrisches Feld?

Wir werden auch bei der Bearbeitung solcher Fragen immer wieder auf die bekannten Rechenregeln aus dem Eindimensionalen zurückkommen, müssen aber einige neue Konzepte kennenlernen, die uns erlauben, geeignete Ableitungen und Integrale zu formulieren.

Überblick über das Kapitel

Wir werden zunächst sogenannte “Kurven” behandeln, das sind Funktionen $f : \mathbb{R} \rightarrow \mathbb{R}^n$, und Problemstellungen wie die Berechnung der Kurvenlänge lösen. Im zweiten Abschnitt wollen wir uns dem umgekehrten Fall zuwenden, Funktionen $f : \mathbb{R}^n \rightarrow \mathbb{R}$, und den sehr wichtigen Begriff der “partiellen Ableitung” kennenlernen. Erst dann wenden wir uns allgemeinen Funktionen $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ zu, verallgemeinern den Begriff der Ableitung $f'(x)$, und werden eine grundlegende Eigenschaft der Ableitung kennenlernen, nämlich, dass sie eine Approximation für f in der Umgebung von x liefert: $f(y) \approx f(x) + f'(x)(y - x)$. Im darauffolgenden vierten Abschnitt behandeln wir einige Aspekte der Integration von Funktionen im Mehrdimensionalen, die für Sie einmal nützlich sein können.

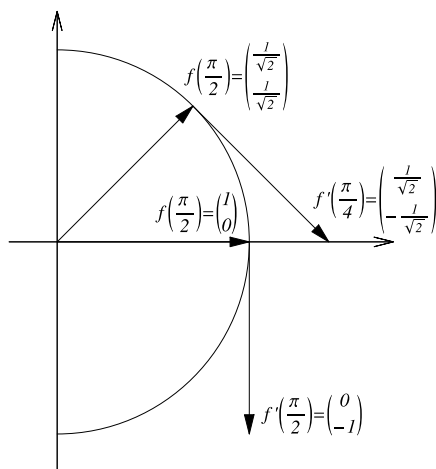


Abbildung 7.1: Halbkreisbogen mit Tangentialvektoren

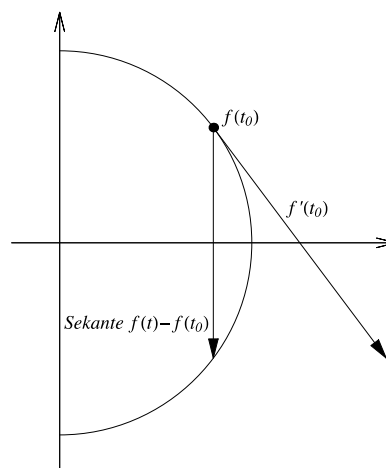


Abbildung 7.2: Der Tangentialvektor $f'(t_0)$ zeigt in die Richtung einer infinitesimal kurzen Sekante $f(t) - f(t_0)$.

7.1 Kurven

Definition 7.1.1 (Kurve).

Eine Kurve im $\mathbb{R}^n$ ist eine stetige Abbildung

$$f : I \rightarrow \mathbb{R}^n, \\ t \mapsto \begin{bmatrix} f_1(t) \\ \vdots \\ f_n(t) \end{bmatrix},$$

wobei $I \subset \mathbb{R}$ ein Intervall ist. Man schreibt $f \in C^k(I, \mathbb{R}^n)$, wenn **alle** Funktionen $f_1, \dots, f_n$ k -mal stetig differenzierbar sind, also $f_1, \dots, f_n \in C^k(I, \mathbb{R})$.

Beispiel 7.1.2 (Halbkreisbogen).

Wir betrachten als erstes Beispiel einer Kurve einen Halbkreisbogen im $\mathbb{R}^2$ (siehe Abbildung 7.1), der durch eine Funktion $f \in C^\infty([0, \pi], \mathbb{R}^2)$ beschrieben wird:

$$f : [0, \pi] \rightarrow \mathbb{R}^2 \\ t \mapsto f(t) = \begin{pmatrix} \sin t \\ \cos t \end{pmatrix}$$

Eine erste Frage, die wir uns stellen können, ist, wie man die Tangente einer Kurve an einer bestimmten Stelle berechnen kann. Dies hängt ganz direkt mit dem Begriff der Ableitung zusammen.

Definition 7.1.3 (Tangentialvektor).

Die **Ableitung** f' einer Kurve $f \in C^1(I, \mathbb{R}^n)$ ist die Abbildung

$$f' : I \rightarrow \mathbb{R}^n$$

$$t \mapsto f'(t) := \begin{bmatrix} f'_1(t) \\ \vdots \\ f'_n(t) \end{bmatrix}.$$

Der Vektor $f'(t)$ heisst **der Tangentialvektor an der Stelle t** (oder auch einfach nur **die Ableitung an der Stelle t**).

Beispiel 7.1.4 (Ableitung des Halbkreisbogens).

Wir berechnen die Ableitung zum Halbkreisbogen von Beispiel 7.1.4:

$$f(t) = \begin{bmatrix} \sin t \\ \cos t \end{bmatrix} \Rightarrow f'(t) = \begin{bmatrix} (\sin t)' \\ (\cos t)' \end{bmatrix} = \begin{pmatrix} \cos t \\ -\sin t \end{pmatrix}.$$

Die Interpretation der Ableitung $f'(t)$ an der Stelle t als Tangentialvektor wird in Abbildung 7.1 für die Stellen $t = \frac{\pi}{4}$ und $t = \frac{\pi}{2}$ veranschaulicht.

Bemerkung 7.1.5. Die Ableitung $f'(t_0)$ kann als Grenzwert

$$f'(t_0) = \lim_{t \rightarrow t_0} \frac{f(t) - f(t_0)}{t - t_0} \quad (7.1)$$

aufgefasst werden, wobei der Vektor $f(t) - f(t_0)$ die Sekante zwischen zwei Punkten der Kurve ist (siehe Abbildung 7.2).

7.1.1 Wie berechnet man die Kurvenlänge?

Eine häufige Fragestellung ist die Berechnung der Länge $L(f)$ einer (gekrümmten) Kurve $f : [t_0, t_f] \rightarrow \mathbb{R}^n$. Die Fragestellung der Berechnung des Kreisumfangs war beispielsweise bereits in der Antike bekannt und stellte für viele Jahrhunderte ein großes Problem dar. Wir können uns heute zum Glück eine einfach auszuwertende Formel zur Berechnung der Länge fast beliebiger Kurven herleiten (also auch der des Kreises), die interessanterweise die Ableitung f' verwendet. Wir starten mit der Beobachtung, dass die Länge einer geraden Strecke, z.B. einer Sekante $f(t_1) - f(t_0)$, ganz einfach durch die Euklidische Vektornorm $\|f(t_1) - f(t_0)\|$ gegeben ist. Die Idee zur Berechnung der Kurvenlänge (oder auch Bogenlänge) $L(f)$ ist nun, die Kurve in kleine Stückchen zu unterteilen, jedes Stück durch die Sekante zwischen seinem Anfangs- und Endpunkt zu ersetzen, und die Gesamtsumme der Sekantenlängen als Approximation der Kurvenlänge $L(f)$ zu verwenden, siehe Figur 7.3.

Um die kleinen Kurvenstücke zu erhalten, wählen wir uns eine große natürliche Zahl N und unterteilen das Intervall $[t_0, t_f]$ in N Stücke $[t_i, t_{i+1}]$, mit

$$t_i = t_0 + i \cdot \frac{t_f - t_0}{N}, \quad i = 0, \dots, N.$$

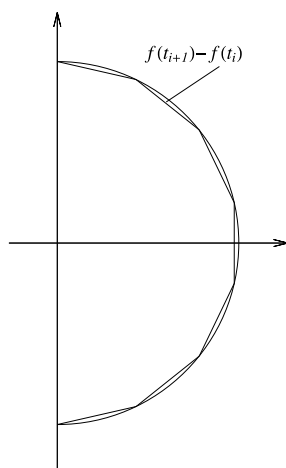


Abbildung 7.3: Approximation der Kurvenlänge durch Sekanten.

Sodann approximieren wir

$$L(f) \approx \sum_{i=0}^{N-1} \|f(t_{i+1}) - f(t_i)\|.$$

Wenn wir nun aber N groß werden lassen, also gleichzeitig auch die Intervalllänge $t_{i+1} - t_i = \frac{t_f - t_0}{N}$ klein, dann erhalten wir eine exakte Darstellung der Kurvenlänge

$$\begin{aligned} L(f) &:= \lim_{N \rightarrow \infty} \sum_{i=0}^{N-1} \|f(t_{i+1}) - f(t_i)\| = \lim_{N \rightarrow \infty} \sum_{i=0}^{N-1} \frac{\|f(t_{i+1}) - f(t_i)\|}{t_{i+1} - t_i} (t_{i+1} - t_i) \\ &= \lim_{N \rightarrow \infty} \sum_{i=0}^{N-1} \left\| \frac{f(t_{i+1}) - f(t_i)}{t_{i+1} - t_i} \right\| (t_{i+1} - t_i) = \lim_{N \rightarrow \infty} \sum_{i=0}^{N-1} \|f'(t_i)\| (t_{i+1} - t_i) \\ &= \int_{t_0}^{t_f} \|f'(t)\| \cdot dt. \end{aligned}$$

In den ersten zwei Schritten haben wir rein algebraische Umformungen verwendet, während wir im dritten Schritt Gleichung (7.1) und im letzten Schritt die Riemann-Definition des Integrals durch Treppensummen ausgenutzt haben. Wir definieren uns also:

Definition 7.1.6 (Kurvenlänge). Die **Länge einer Kurve** $f \in C^1([t_0, t_f], \mathbb{R}^n)$ ist

$$L(f) = \int_{t_0}^{t_f} \|f'(t)\| \, dt.$$

Beispiel 7.1.7 (Halbkreisumfang).

Wir betrachten wieder den Halbkreisbogen aus Beispielen 7.1.2 und 7.1.4, und berechnen seine Kurvenlänge $L(f)$.

$$\begin{aligned} f'(t) &= \begin{pmatrix} \cos t \\ -\sin t \end{pmatrix}, \\ \|f'(t)\| &= \sqrt{f_1^2(t) + f_2^2(t)} = \sqrt{(\cos t)^2 + (\sin t)^2} = 1, \\ L(f) &= \int_0^\pi \|f'(t)\| dt = \int_0^\pi 1 \cdot dt = \pi. \end{aligned}$$

Beispiel 7.1.8 (Einfach-Helix im $\mathbb{R}^3$).

Wir betrachten als ein dreidimensionales Beispiel die folgende Helix, die in Abbildung 7.4 gezeigt ist:

$$\begin{aligned} f(t) &= f : [0, 4\pi] \rightarrow \mathbb{R}^3, f(t) = \begin{pmatrix} t \\ \cos t \\ \sin t \end{pmatrix}, \\ f'(t) &= \begin{pmatrix} 1 \\ -\sin t \\ \cos t \end{pmatrix}, \\ \|f'(t)\| &= \sqrt{(1)^2 + (-\sin t)^2 + (\cos t)^2} = \sqrt{1 + 1} = \sqrt{2}, \\ L(t) &= \int_0^{4\pi} \sqrt{2} \, dt = 4\sqrt{2}\pi. \end{aligned}$$

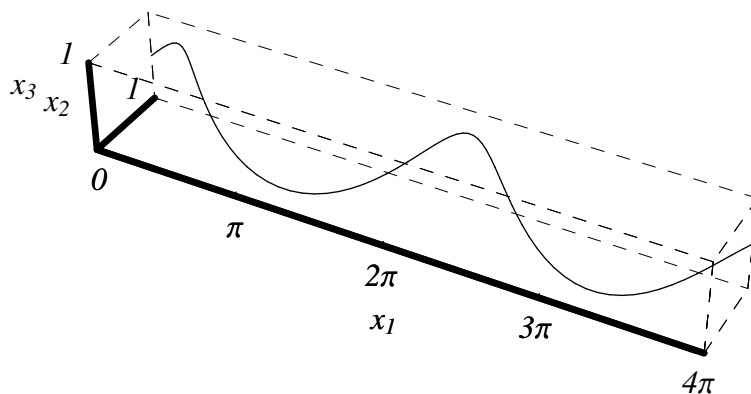


Abbildung 7.4: Helix im $\mathbb{R}^3$

Beispiel 7.1.9 (Seltsame Spirale im $\mathbb{R}^2$).

Als ein weiteres Beispiel betrachten wir die folgende, etwas seltsam anmutende Spiralkurve im $\mathbb{R}^2$, und geben einen Ausdruck für Ihre Kurvenlänge an.

$$\begin{aligned} f(t) &= f : [0, 2\pi] \rightarrow \mathbb{R}^2, f(t) = \begin{pmatrix} t + 2 \sin t \\ \cos t \end{pmatrix}, \\ f'(t) &= \begin{pmatrix} 1 + 2 \cos t \\ -\sin t \end{pmatrix}, \\ \|f'(t)\| &= \sqrt{(1 + 2 \cos t)^2 + (-\sin t)^2}, \\ L(t) &= \int_0^{2\pi} \sqrt{(1 + 2 \cos t)^2 + \sin^2 t} \, dt. \end{aligned}$$

Diesen Ausdruck könnten wir nun durch geeignete Integralumformungen oder mit Hilfe des Computers berechnen.

7.2 Ableitungen im $\mathbb{R}^n$

Im Gegensatz zum vorherigen Abschnitt wollen wir nun den Fall von Funktionen betrachten, die nicht nur von einem Argument abhängen (das wir t genannt hatten), sondern gleich von mehreren, die wir hier meist mit $x_1, \dots, x_n$ bezeichnen. Der Einfachheit halber betrachten wir zunächst nur skalare Funktionen $f : \mathbb{R}^n \rightarrow \mathbb{R}$.

7.2.1 Veranschaulichung von Funktionen mehrerer Argumente

Zunächst stellen wir uns die Frage, wie man sich solche Funktionen veranschaulichen kann. Man kann dies im Wesentlichen auf zwei verschiedene Weisen machen, die wir für den Fall $n = 2$, der uns am meisten interessiert, ganz intuitiv erfassen.

Möglichkeit 1: Veranschaulichung als Graph

Wir betrachten eine Art Gebirgslandschaft, nämlich den *Graph* der Funktion im $\mathbb{R}^{n+1}$, also die Menge der Punkte

$$\{(x_1, \dots, x_n, y) \in \mathbb{R}^{n+1} \mid y = f(x_1, \dots, x_n)\},$$

der in Abbildung 7.5 illustriert ist.

Möglichkeit 2: Veranschaulichung durch Niveaumengen

Die zweite Möglichkeit der Veranschaulichung ist den Kartographen von Gebirgslandschaften nachempfunden, die einfach Höhenlinien auf Karten einzeichnen. Mathematisch exakt nennen wir diese Höhenlinien jetzt *Niveaumengen*.

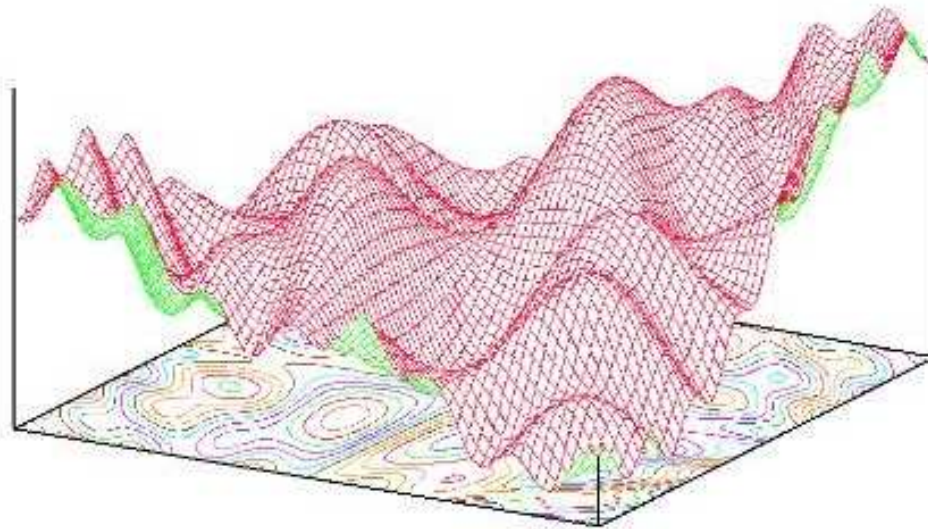


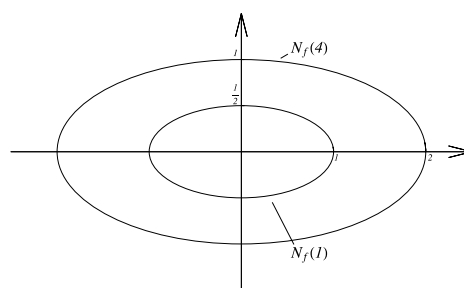
Abbildung 7.5: Zusammenhang zwischen Graph und Niveaulinien.

Definition 7.2.1 (Niveaumenge).

Die **Niveaumenge** $N_f(c)$ einer Funktion $f : U \rightarrow \mathbb{R}$ zum Wert c ist definiert als das Urbild von c unter f , also als die Menge

$$N_f(c) := \{x \in U \mid f(x) = c\}.$$

In Abbildung 7.5 ist der Zusammenhang zwischen Graph und Niveaumengen illustriert.

Abbildung 7.6: Niveaumengen der Funktion $f(x) = x_1^2 + 4x_2^2$.

Beispiel 7.2.2. Wir zeigen in Abbildung 7.6 zwei Niveaumengen für die Funktion

$$\begin{aligned} f : \mathbb{R}^2 &\rightarrow \mathbb{R}, \\ x &\mapsto x_1^2 + 4x_2^2. \end{aligned}$$

7.2.2 *Offene Mengen

Um für Funktionen mehrerer Argumente mathematisch korrekt den Begriff der Ableitung definieren zu können, führen wir zunächst einen abstrakten Begriff aus dem Gebiet der Topologie ein, den Begriff der *offenen Menge*.

Definition 7.2.3 (Offene Menge).

Eine Menge $U \subset \mathbb{R}^n$ heisst **offen**, falls für jedes $x \in U$ ein $\epsilon > 0$ existiert, so dass der ϵ -Ball

$$B(x, \epsilon) := \{y \in \mathbb{R}^n \mid \|y - x\| \leq \epsilon\}$$

ganz in U enthalten ist, also $B(x, \epsilon) \subset U$.

Beispiel 7.2.4.

1. $U = \{x \in \mathbb{R}^2 \mid x_1 > 0, x_2 > 0\}$ ist offen, denn mit $\epsilon = \frac{1}{2}\min(x_1, x_2)$ gilt $B(x, \epsilon) \subset U$, siehe Abbildung 7.7.

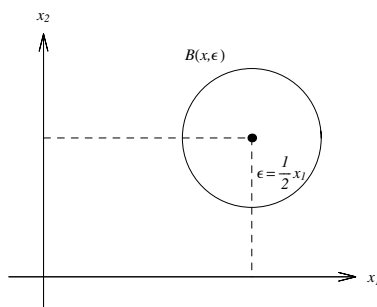


Abbildung 7.7: Beispiel 7.2.4.1: Für jedes $(x_1, x_2) \in U$ gibt es einen ϵ -Ball, der ganz in U enthalten ist.

2. $U = \{x \in \mathbb{R}^2 \mid x_1 \geq 0, x_2 > 0\}$ ist nicht offen, denn für $(x_1, x_2) = (0, 1) \in U$ gibt es keinen ϵ -Ball, der ganz in U enthalten ist.
3. $U = \{x \in \mathbb{R}^2 \mid x_1 = 0, x_2 > 0\}$ ist aus dem gleichen Grund nicht offen.

Wir betrachten im folgenden immer eine Funktion

$$f : U \rightarrow \mathbb{R}, \quad (x_1, \dots, x_n) \mapsto f(x_1, \dots, x_n)$$

wobei $U \subset \mathbb{R}^n$ *offen* ist. Diese Annahme stellt sicher, dass es für jeden Punkt aus $x_0 \in U$ eine ganze ϵ -Umgebung gibt, für den die Funktion definiert ist, und dies erlaubt uns, die folgenden Grenzwerte mathematisch korrekt zu definieren.

7.3 Allgemeines Konzept der Ableitung

Eine Ableitung oder Differentiation ist eine lokale, lineare Approximation nichtlinearer Abbildungen.

Wiederholung der Definition einer linearen Abbildung

$$f(u + v) = f(u) + f(v)$$

$$f(\lambda u) = \lambda f(u)$$

Bezeichnung: Vektorraum = linearer Raum

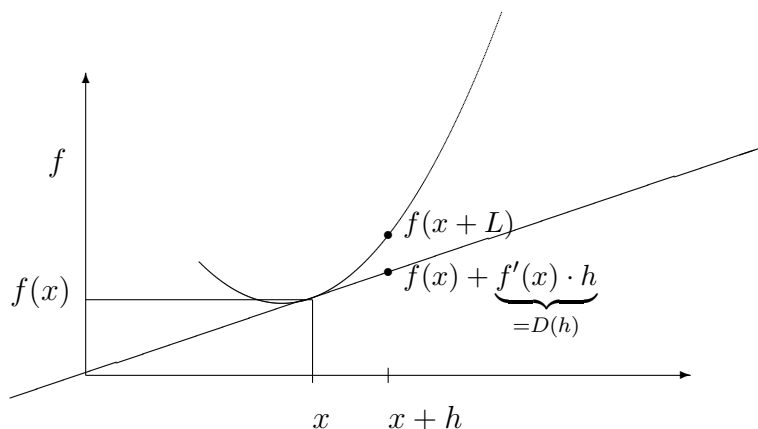
Eine Taylor-Entwicklung erster Ordnung von f liefert

$$f(x + h) \approx f(x) + f'(x) \cdot (x + h - x) = f(x) + f'(x) \cdot h \quad (7.2)$$

„für kleine h “.

Die Abbildung $D : h \mapsto f'(x) \cdot h$ (7.2) ist **linear** bezüglich h . Sie beschreibt die Geradengleichung der Tangente an den Graphen von f im Punkt x .

Lineare lokale Approximation



Abstrakte Definition einer Ableitung als lineare, lokale Approximation möglich! Damit kann man die Operation „Differentiation“ für Abbildungen f zwischen beliebigen normierten Vektorräumen V und W definieren.

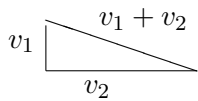
Wiederholung der Norm

Definition 7.3.1. Eine Abbildung $\| \cdot \| : V \rightarrow \mathbb{R}$ mit Eigenschaften

- a) $\|v\| = 0 \Leftrightarrow v = 0$
- b) $\|\lambda v\| = |\lambda| \cdot \|v\| \quad \forall \lambda \in \mathbb{R}, v \in V$
- c) $\|v_1 + v_2\| \leq \|v_1\| + \|v_2\|$ (Dreiecksungleichung)

heißt Norm auf V .

Veranschaulichung von c) (Dreiecksungleichung)



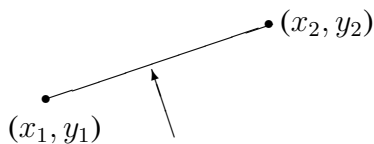
Die Norm ist ein verallgemeinerter „Abstandsbegriff“.

Bemerkung 7.3.2.

- Es gibt in der Regel verschiedene Möglichkeiten, auf einem Vektorraum V eine Norm $\|\cdot\|$ zu definieren.
- Die gängigste Norm ist die sogenannte Euklidische Norm, die von einem Skalarprodukt induziert wird: $\|u\| := \sqrt{\langle u, u \rangle}$. Im $\mathbb{R}^n$ ist das

$$\|u\| := \sqrt{\begin{pmatrix} u_1 \\ \vdots \\ u_n \end{pmatrix} \cdot \begin{pmatrix} u_1 \\ \vdots \\ u_n \end{pmatrix}} = \sqrt{\sum_{i=1}^n u_i^2} \quad (7.3)$$

Der Abstand zweier Punkte $(x_1, y_1) \in \mathbb{R}^2$ und $(x_2, y_2) \in \mathbb{R}^2$ ist die Länge der Verbindungsstrecke in der Euklidischen Norm:



$$\left\| \begin{pmatrix} x_1 - x_2 \\ y_1 - y_2 \end{pmatrix} \right\| = \sqrt{(x_1 - x_2)^2 + (y_1 - y_2)^2}$$

(Pythagoras)

7.4 Konvergenz und Stetigkeit

Bemerkung 7.4.1. Die Konzepte im vorherigen Abschnitt wurden hier speziell für den $\mathbb{R}^n$ eingeführt, gelten aber im Allgemeinen analog für beliebige normierte Räume.

Der Ableitungsbegriff basiert auf der Konvergenz !

Definition 7.4.2 (Konvergenz). Sei $(x^k)_{k \in \mathbb{N}}$ eine Folge im $\mathbb{R}^n$, d.h. die Folgenglieder sind „Punkte“ bzw. „Vektoren“ $(x_1^k, x_2^k, \dots, x_n^k)^T$ [k ist hier ein Index, keine Potenz !].

Dann heißt die Folge konvergent gegen $a \in \mathbb{R}^n$ $\left(\lim_{k \rightarrow \infty} x^k = a \right)$, falls gilt:

$\forall \epsilon > 0 \exists N \in \mathbb{N}$ mit

$$x^k \in B(a, \epsilon) \quad \forall k \geq N$$

$$\text{d.h. } \|x^k - a\| \leq \epsilon \quad \forall k \geq N$$

Dabei ist $B(a, \epsilon)$ der ϵ -Ball um x (siehe Abschnitt ??). Man beachte die Analogie zum Fall $n = 1$ (ϵ -Umgebung) !

Bemerkung 7.4.3. Konvergenz gilt für die jeweils gewählte Norm.

Beispiel 7.4.4. Sei $(x^k)_{k \in \mathbb{N}}$ eine Folge mit $x^k \in \mathbb{R}^2$, $x^k := \begin{pmatrix} 1 + \frac{1}{2^k} \\ \frac{1}{k} \end{pmatrix}$.

Zeige $\lim_{k \rightarrow \infty} x^k = \begin{pmatrix} 1 \\ 0 \end{pmatrix} \in \mathbb{R}^2$!

$$\left\| \begin{pmatrix} 1 + \frac{1}{2^k} \\ \frac{1}{k} \end{pmatrix} - \begin{pmatrix} 1 \\ 0 \end{pmatrix} \right\| = \left\| \begin{pmatrix} \frac{1}{2^k} \\ \frac{1}{k} \end{pmatrix} \right\| = \sqrt{\left(\frac{1}{2^k}\right)^2 + \left(\frac{1}{k}\right)^2} = \sqrt{\underbrace{\frac{1}{2^{2k}}}_{\rightarrow 0} + \underbrace{\frac{1}{k^2}}_{\rightarrow 0}} \leq \epsilon$$

$\forall k \geq N \in \mathbb{N}$.

Für $k \rightarrow \infty$ gilt dann $\frac{1}{2^{2k}} \rightarrow 0$ und $\frac{1}{k^2} \rightarrow 0$ und damit insgesamt Konvergenz.

Definition 7.4.5 (Stetigkeit). Eine Abbildung $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ heißt stetig im Punkt $a \in \mathbb{R}^n$, falls $\lim_{x \rightarrow a} f(x) = f(a)$, d.h. für jede Folge $(x^k)_{k \in \mathbb{N}}$, $x^k \in \mathbb{R}^n$ mit $\lim_{k \rightarrow \infty} x^k = a$ gilt $\lim_{k \rightarrow \infty} f(x^k) = f(a)$. (wieder analog zum Fall von Funktionen $f : \mathbb{R} \rightarrow \mathbb{R}$)

Definition 7.4.6 (Stetigkeit, ϵ - δ -Kriterium). Eine Abbildung $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ ist genau dann stetig im Punkt $a \in \mathbb{R}^n$, wenn gilt:

$$\forall \epsilon > 0 \exists \delta > 0 : \|f(x) - f(a)\| < \epsilon,$$

$\forall x \in \mathbb{R}^n$ mit $\|x - a\| < \delta$. (Analogie zu Funktionen einer Variablen $f : \mathbb{R} \rightarrow \mathbb{R}$ mit der Norm $|| \cdot ||$ anstelle des Betrages $| \cdot |$)

7.4.1 Partielle Ableitungen

Jetzt können wir endlich den Begriff der Ableitung auf Funktionen mit mehreren Argumenten verallgemeinern. Die gewöhnliche Ableitung für $f : \mathbb{R} \rightarrow \mathbb{R}$ kann als Steigung der Funktion aufgefasst werden. Im Falle mehrerer Argumente müssen wir uns fragen, in welcher Richtung wir die Steigung angeben wollen. Die Steigung in einer Koordinatenrichtung, z.B. in Richtung von x_k , nennt man dann einfach die *partielle Ableitung*. Sie ist wie folgt definiert.

Definition 7.4.7 (Partielle Ableitung).

Sei $U \subset \mathbb{R}^n$ offen, und $f : U \rightarrow \mathbb{R}$. Die **partielle Ableitung von f nach x_k** , falls sie existiert, ist die Funktion

$$\begin{aligned} \frac{\partial f}{\partial x_k} : U &\rightarrow \mathbb{R}, \\ x &\mapsto \frac{\partial f}{\partial x_k}(x), \end{aligned}$$

wobei $\frac{\partial f}{\partial x_k}(x)$ als der Limes

$$\frac{\partial f}{\partial x_k}(x_1, \dots, x_n) := \lim_{h \rightarrow 0} \frac{f(x_1, \dots, x_k + h, \dots, x_n) - f(x_1, \dots, x_k, \dots, x_n)}{h}$$

definiert ist.

Man bezeichnet $\frac{\partial f}{\partial x_k}$ manchmal auch als **k-te partielle Ableitung von f** , oder man schreibt statt $\frac{\partial f}{\partial x_k}$ kurz $\partial_k f$.

Der obige Limes sieht komplizierter aus, als er ist, denn er besagt nichts anderes, als dass die Funktion beim Ableiten einfach nur als Funktion des einen Argumentes x_k betrachtet werden soll, also alle anderen $n - 1$ Argumente beim Ableiten als *Konstante* angesehen werden. Daraus folgt ganz natürlich die Rechenregel zur Berechnung partieller Ableitungen, die am besten anhand von Beispielen deutlich wird.

Beispiel 7.4.8 (für partielle Ableitungen).

1. Die Funktion $f(x_1, x_2) = x_1^2 \cdot x_2^4 + x_2$ hat die partiellen Ableitungen:

$$\frac{\partial f}{\partial x_1}(x_1, x_2) = 2x_1 \cdot x_2^4 + 0 \quad (x_2 \text{ wird als Konstante behandelt})$$

$$\frac{\partial f}{\partial x_2}(x_1, x_2) = x_1^2 \cdot 4x_2^3 + 1 \quad (x_1 \text{ wird als Konstante behandelt})$$

2. Die Funktion $f(x_1, x_2, x_3) = \sin(x_1) \cdot (x_1 + 3x_2x_3)$ hat die partiellen Ableitungen:

$$\frac{\partial f}{\partial x_1}(x_1, x_2, x_3) = \cos(x_1) \cdot (x_1 + 3x_2x_3) + \sin(x_1) \cdot (1 + 0)$$

$$\frac{\partial f}{\partial x_2}(x_1, x_2, x_3) = \sin(x_1) \cdot (0 + 3x_3)$$

$$\frac{\partial f}{\partial x_3}(x_1, x_2, x_3) = \sin(x_1) \cdot (0 + 3x_2)$$

7.4.2 Totale Ableitung

Wenn man alle partiellen Ableitungen in einem Vektor zusammenfasst, erhält man, ähnlich wie zuvor bei den Kurven, einen Vektor, den wir die *totale Ableitung* nennen. Es gilt folgende Definition.

Definition 7.4.9 (Stetige Differenzierbarkeit, Totale Ableitung).

Eine Funktion $f : U \rightarrow \mathbb{R}$ ($U \subset \mathbb{R}^n$ offen) heisst **stetig differenzierbar**, $f \in C^1(U, \mathbb{R})$, wenn für jedes $x \in U$ **alle** Ableitungen $\frac{\partial f}{\partial x_1}(x) \dots \frac{\partial f}{\partial x_n}(x)$ existieren und stetig sind.

Für stetig differenzierbare Funktionen $f \in C^1(U, \mathbb{R})$ heisst der Zeilenvektor

$$f'(x) = \left(\frac{\partial f}{\partial x_1}(x), \dots, \frac{\partial f}{\partial x_n}(x) \right)$$

die **(totale) Ableitung von f an der Stelle x** .

Der Grund dafür, dass wir den Vektor $f'(x)$ totale Ableitung nennen, wird durch den folgenden sehr wichtigen Satz deutlich, der eine wichtigste Eigenschaft der “normalen” Ableitung von $f : \mathbb{R} \rightarrow \mathbb{R}$ verallgemeinert, für die nämlich gilt

$$f'(\bar{x}) = \lim_{x \rightarrow \bar{x}} \frac{f(x) - f(\bar{x})}{x - \bar{x}} \Leftrightarrow \lim_{x \rightarrow \bar{x}} \frac{f(x) - (f(\bar{x}) + f'(\bar{x})(x - \bar{x}))}{x - \bar{x}} = 0,$$

und die wir so interpretieren können, dass für festes $\bar{x}$ der Ausdruck $f(\bar{x}) + f'(\bar{x})(x - \bar{x})$ eine Approximation für $f(x)$ ist, wobei der Fehler $\phi(x, \bar{x}) := f(x) - (f(\bar{x}) + f'(\bar{x})(x - \bar{x}))$ für $x \rightarrow \bar{x}$ *schneller* gegen Null konvergiert als $x - \bar{x}$. Für eine Funktion $f : \mathbb{R}^n \rightarrow \mathbb{R}$, bei der die totale Ableitung $f'(\bar{x})$ ein Zeilenvektor ist, gilt nun analog:

Satz 7.4.10 (Approximationseigenschaft der Ableitung).

Sei $U \subset \mathbb{R}^n$ offen und $f \in C^1(U, \mathbb{R})$. Dann gilt für alle $x, \bar{x} \in U$ dass

$$f(x) = f(\bar{x}) + f'(\bar{x})(x - \bar{x}) + \underbrace{\phi(\bar{x}, x)}_{\text{Fehler}}, \quad \text{wobei} \quad \lim_{x \rightarrow \bar{x}} \frac{\phi(\bar{x}, x)}{\|x - \bar{x}\|} = 0.$$

Wir schreiben oft auch einfach nur

$$f(x) \approx f(\bar{x}) + f'(\bar{x})(x - \bar{x})$$

und sagen, $f(\bar{x}) + f'(\bar{x})(x - \bar{x})$ ist eine Approximation **erster Ordnung** für $f(x)$. Man beachte, dass $(x - \bar{x})$ ein stehender Vektor ist, und der liegende Vektor $f'(\bar{x})$ als lineare Abbildung $f'(x_0) : \mathbb{R}^n \rightarrow \mathbb{R}$ aufgefasst werden kann.

***Bemerkung 7.4.11** (Verallgemeinerte Ableitung).

Wir bemerken noch für mathematisch Interessierte, dass die im Satz gegebene Approximationseigenschaft der Ableitung $f'(x_0)$ oft sogar zur *Definition* der Ableitung verwendet wird. Dies erlaubt eine ganz weitgehende Verallgemeinerung der Ableitung auf alle Abbildungen zwischen zwei Vektorräumen, die jeweils mit einer Norm ausgestattet sind.

***Bemerkung 7.4.12** (Stetigkeit).

Als eine zweite Bemerkung für mathematisch Interessierte erwähnen wir, dass die bloße *Existenz* der partiellen Ableitungen nicht ausreicht, um diese Approximationseigenschaft zu beweisen, sondern dass wir voraussetzen müssen, dass die partiellen Ableitungen auch *stetig* sind. Deshalb haben wir in der Definition der totalen Ableitung in Definition 7.4.9 die *stetige* Ableitbarkeit vorausgesetzt.

Wir wollen veranschaulichen, was die Approximationseigenschaft der Ableitung wirklich bedeutet. Dafür denken wir uns $\bar{x}$ als einen gegebenen festen Vektor, und definieren den Abweichungsvektor $\Delta x := x - \bar{x}$ sowie die Abweichung $\Delta f := f(x) - f(\bar{x})$. Betrachten wir also

$$\begin{aligned} f(x) \approx f(\bar{x}) + f'(\bar{x}) \cdot (x - \bar{x}) &\Leftrightarrow f(x) - f(\bar{x}) \approx f'(\bar{x}) \cdot (x - \bar{x}) \\ &\Leftrightarrow \Delta f \approx f' \cdot \Delta x \\ &\Leftrightarrow \Delta f \approx \frac{\partial f}{\partial x_1} \cdot \Delta x_1 + \frac{\partial f}{\partial x_2} \cdot \Delta x_2 + \cdots + \frac{\partial f}{\partial x_n} \Delta x_n \end{aligned} \quad (7.4)$$

wobei jetzt alle Terme Skalare sind. Man könnte also sagen: „Die Änderung von f ist in erster Ordnung eine gewichtete Summe der Änderungen der Argumente $x_1, \dots, x_n$.“

Definition 7.4.13 (Totales Differential).

Der Ausdruck

$$df := \frac{\partial f}{\partial x_1}(x_1, \dots, x_n) dx_1 + \cdots + \frac{\partial f}{\partial x_n}(x_1, \dots, x_n) dx_n$$

wird **totales Differential** genannt.

Motivation für diese Definition: Wie ändert sich f , wenn man alle x_i um jeweils Δx_i verändert?

$$\Delta f \approx \frac{\partial f}{\partial x_1}(x_1, \dots, x_n) \Delta x_1 + \cdots + \frac{\partial f}{\partial x_n}(x_1, \dots, x_n) \Delta x_n = \sum_{i=1}^n \frac{\partial f}{\partial x_i} \Delta x_i$$

nach Grenzübergang in Leibniz-Notation ($\Delta \rightarrow 0$):

$$df = \sum_{i=1}^n \frac{\partial f}{\partial x_i}(x_1, \dots, x_n) dx_i, \quad (7.5)$$

wo $\frac{\partial f}{\partial x_i}(x_1, \dots, x_n) dx_i$ das Skalarprodukt $f'(x) \cdot \begin{pmatrix} dx_1 \\ \vdots \\ dx_n \end{pmatrix}$ ist.

Diesen Ausdruck nennt man auch „**totales Differential**“ von f .

Totale Differentiale in Leibniz-Notation werden häufig in der Thermodynamik (Physikalische Chemie, Theoretische Physik) benutzt.

Tangentialebene und totales Differential

Die totale Ableitung entspricht einer lokalen, linearen Approximation von f . Für $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ ist das die sogenannte Tangentialebene.

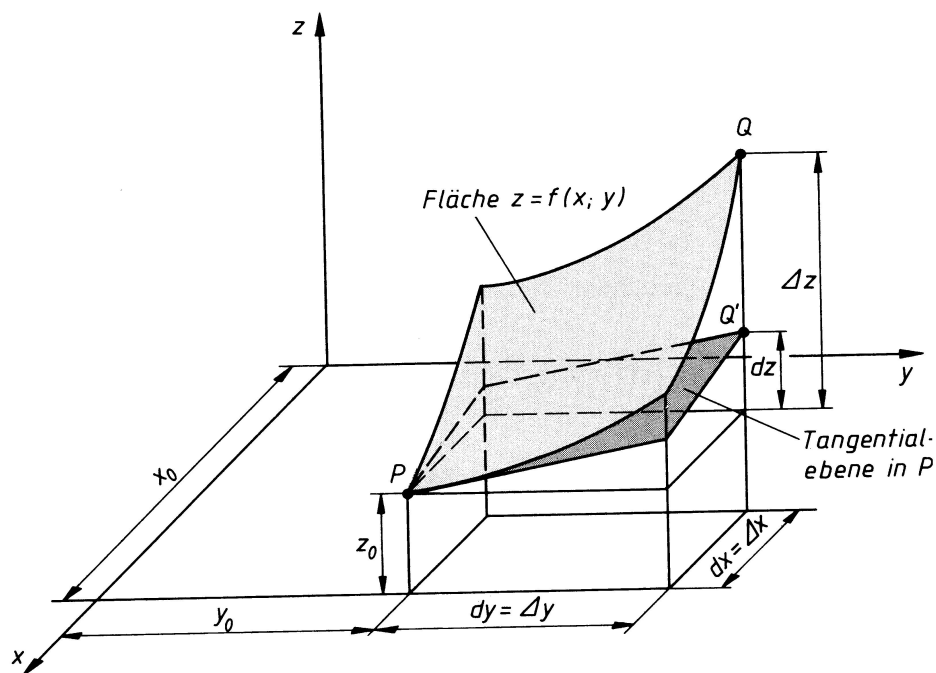


Abbildung 7.8: Abbildung aus [Pap]

Beispiel 7.4.14 (Ideales Gasgesetz).

Wir betrachten als Funktion die Abhängigkeit des Drucks p von den zwei Variablen T = Temperatur und V = Volumen. Nach dem idealen Gasgesetz ergibt sich:

$$p(T, V) = \frac{n \cdot k \cdot T}{V},$$

wobei n = Zahl der Moleküle und k die Boltzmannkonstante ist. Wie ändert sich der Druck bei kleineren Änderungen der Temperatur und/oder des Volumens? Wir berechnen das totale

Differential

$$\begin{aligned}
 dp &= \frac{\partial p}{\partial T}(T, V) dT + \frac{\partial p}{\partial V}(T, V) dV \\
 &= \frac{nk}{V} \cdot dT + \left(-\frac{nkT}{V^2} \right) \cdot dV \\
 &= \frac{nk}{V} \left(dT - \frac{T}{V} dV \right)
 \end{aligned}$$

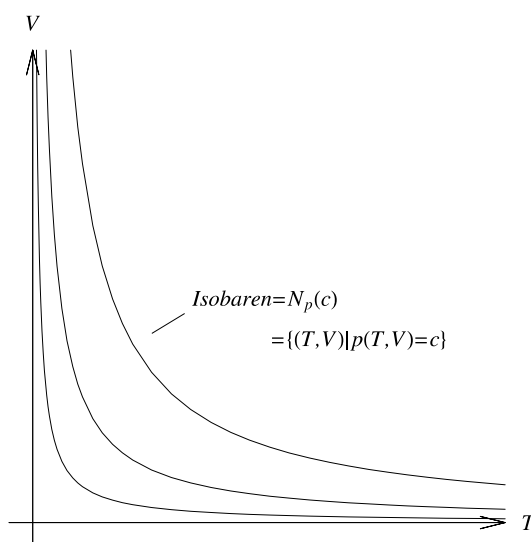


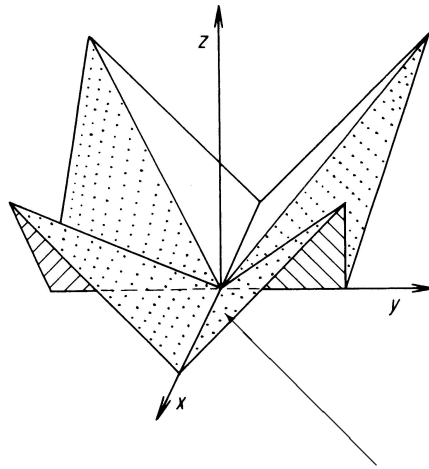
Abbildung 7.9: Isobaren für das ideale Gasgesetz.

Daraus sehen wir sofort:

- Bei Temperaturerhöhung steigt der Druck, denn $\frac{nk}{V}$ ist positiv.
- Bei Volumenvergrößerung sinkt der Druck, denn $-\frac{nkT}{V^2}$ ist negativ.
- Falls Temperaturerhöhung dT und Volumenvergrößerung dV in dem Zusammenhang $dT = \frac{T}{V}dV$ stehen, bleibt der Druck konstant.

Der letzte Punkt erlaubt einem z.B. direkt die Steigung $dV/dT = V/T$ einer Isobaren zu ermitteln, also einer Linie in der T, V -Ebene, auf der der Druck konstant ist, wie in Abbildung 7.9 illustriert. Das Rechnen mit totalen Differentialen ist häufig sehr praktisch.

Auch wenn $\frac{\partial f}{\partial x}(0, 0)$, $\frac{\partial f}{\partial y}(0, 0)$ existieren, muss nicht notwendig eine Tangentialebene existieren ! (siehe Abb. 7.4.2).



Stetigkeit von $\frac{\partial f}{\partial x}(0,0)$ und $\frac{\partial f}{\partial y}(0,0)$ verletzt!

Abbildung 7.10: Nichtexistenz einer Tangentialebene, Abb. aus [Pap]

7.4.3 Vollständiges Differential

Definition 7.4.15. Ein Ausdruck der Form $g(x_1, x_2)dx_1 + h(x_1, x_2)dx_2$ heißt vollständiges (exaktes, totales) Differential, wenn es eine Funktion $f(x_1, x_2)$ gibt, so dass

$$df = \frac{\partial f}{\partial x_1}dx_1 + \frac{\partial f}{\partial x_2}dx_2 = g(x_1, x_2)dx_1 + h(x_1, x_2)dx_2$$

gilt, d.h. in Kurzform

$$g = \frac{\partial f}{\partial x_1}, h = \frac{\partial f}{\partial x_2}$$

Hinreichende Bedingung für Vollständigkeit des Differentials

Satz 7.4.16. Seien $g, h : U \rightarrow \mathbb{R}$ stetig differenzierbar und $U \subset \mathbb{R}^2$ einfach zusammenhängend. Dann gilt:

$$\frac{\partial g}{\partial x_2} = \frac{\partial h}{\partial x_1} \Rightarrow g(x_1, x_2)dx_1 + h(x_1, x_2)dx_2$$

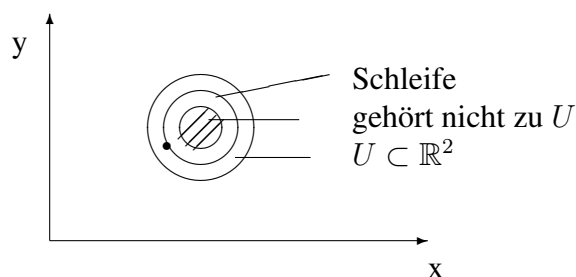
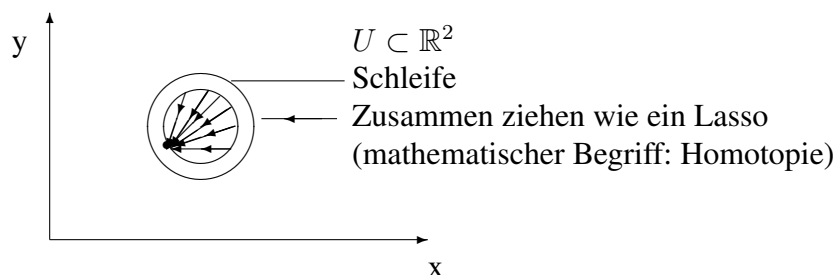
ist vollständiges (exaktes) Differential.

Exkurs: Erläuterung „einfach zusammenhängend“

Im anschaulichen Sinne meint „einfach zusammenhängend“ für eine Teilmenge vom $\mathbb{R}^2$, dass jede Schleife in U stetig auf einen Punkt zusammengezogen werden kann.

U ist in diesem Beispiel einfach zusammenhängend. Schleife ist hier nicht stetig auf einen Punkt zusammenziehbar. $\rightarrow U$ nicht einfach zusammenhängend.

In Anwendungen: U fast immer einfach zusammenhängend ($\rightarrow$ Satz 7.5.2 anwendbar!)



Beispiel 7.4.17. Zeige, dass das Differential $2xy^3dx + 3x^2y^2dy$ vollständig ist, d.h. dass es eine Funktion $f(x, y)$ gibt, so dass $df = \frac{\partial f}{\partial x}dx + \frac{\partial f}{\partial y}dy = 2xy^3dx + 3x^2y^2dy$ gilt.
Überprüfe hinreichende Bedingung

$$\frac{\partial(2xy^3)}{\partial y} = \frac{\partial(3x^2y^2)}{\partial x} ?$$

(hier ist f durch „Hinsehen“ bestimmbar: $f = x^2y^3$)

$$6xy^2 = 6xy^2$$

$\Rightarrow$ Differential vollständig.

Wichtige Interpretation

Die Funktion f mit der Eigenschaft $df = gdx_1 + hdx_2$, die im Falle eines vollständigen Differentials existiert, ist eine Verallgemeinerung einer Stammfunktion für Abbildungen $\mathbb{R} \rightarrow \mathbb{R}$ auf Abb. $\mathbb{R}^2 \rightarrow \mathbb{R}$. Weitere Verallgemeinerung auf Abb. $\mathbb{R}^n \rightarrow \mathbb{R}$ möglich: Differential $df = \sum_{i=1}^n \frac{\partial f}{\partial x_i} dx_i \rightarrow$ „Stammfunktion“ von df ist f (später damit Berechnung von Integralen !)

Wir werden zum Schluss des Abschnitts noch eine interessante zweite Interpretation der totalen Ableitung als „Gradient“ kennenlernen.

Definition 7.4.18 (Gradient).

Sei $f \in C^1(U, \mathbb{R})$ ($U \subset \mathbb{R}^n$ offen), $x \in U$. Den transponierten Ableitungsvektor

$$\nabla f(x) := f'(x)^T = \begin{pmatrix} \frac{\partial f}{\partial x_1}(x) \\ \vdots \\ \frac{\partial f}{\partial x_n}(x) \end{pmatrix}$$

nennt man den **Gradient** von f an der Stelle x (Das Symbol ∇ liest man „Nabla“).

Den Gradient $\nabla f(x)$ ist ein Vektor im $\mathbb{R}^n$ mit den folgenden Eigenschaften:

- $\nabla f(x)$ zeigt in die Richtung **steilsten Anstiegs**, d.h. wenn man im $\mathbb{R}^n$ eine Längeneinheit in Richtung des Gradienten geht, steigt f stärker an als in jeder anderen Richtung.
- Wenn x in einer Niveaumenge $N_f(c)$ liegt, dann steht der Gradient $\nabla f(x)$ orthogonal auf der Niveaumenge $N_f(c)$.
- Wenn $\begin{pmatrix} x \\ f(x) \end{pmatrix} \in \mathbb{R}^{n+1}$ im Graphen von f liegt, dann steht der Vektor $\begin{pmatrix} -\nabla f(x) \\ 1 \end{pmatrix} \in \mathbb{R}^{n+1}$ orthogonal auf dem Graphen.

7.4.4 Partielle Ableitungen höherer Ordnung

Wir können uns fragen, ob man eine Funktion auch mehrmals partiell ableiten kann, also Ableitungen höherer Ordnung bilden kann. Dies geht tatsächlich, wenn die Funktion f ausreichend glatt ist. Man leitet dann eine partielle Ableitung ganz einfach noch ein weiteres Mal ab, indem man z.B. Ausdrücke der Form $\frac{\partial}{\partial x_k} \left(\frac{\partial f}{\partial x_l} \right)(x)$ berechnet.

Definition 7.4.19 (Zweite Partielle Ableitungen).

Sei $f \in C^1(U, \mathbb{R})$, $U \subset \mathbb{R}^n$ offen, und $k, l \in \{1, \dots, n\}$. Den Ausdruck

$$\frac{\partial^2 f}{\partial x_k \partial x_l}(x) := \frac{\partial}{\partial x_k} \left(\frac{\partial f}{\partial x_l} \right)(x)$$

nennt man die zweite partielle Ableitung von f nach x_k und x_l . Existieren alle zweiten partiellen Ableitungen und sind sie stetig, so schreibt man $f \in C^2(U, \mathbb{R})$.

Die Reihenfolge der partiellen Ableitungen ist dann interessanterweise egal:

Satz 7.4.20 (Vertauschung partieller Ableitungen).

Sei $f \in C^2(U, \mathbb{R})$, $U \subset \mathbb{R}^n$ offen, und $k, l \in \{1, \dots, n\}$. Dann gilt:

$$\frac{\partial}{\partial x_k} \left(\frac{\partial f}{\partial x_l} \right)(x) = \frac{\partial}{\partial x_l} \left(\frac{\partial f}{\partial x_k} \right)(x) =: \frac{\partial^2 f}{\partial x_k \partial x_l}(x).$$

Beispiel 7.4.21.

$$\begin{aligned}\frac{\partial^2(x^3y^2)}{\partial x \partial y} &= \frac{\partial}{\partial x} \left(\frac{\partial(x^3y^2)}{\partial y} \right) = \frac{\partial}{\partial x} (x^3 2y) = 6x^2y \\ \text{bzw.} &= \frac{\partial}{\partial y} \left(\frac{\partial(x^3y^2)}{\partial x} \right) = \frac{\partial}{\partial y} (3x^2y^2) = 6x^2y.\end{aligned}$$

7.5 Taylor-Entwicklung und lokale Extrema

Satz 7.5.1. Sei $U \subset \mathbb{R}^n$ offen, $f : U \rightarrow \mathbb{R}$ zweimal stetig differenzierbar, $h \in \mathbb{R}^n$, dann gilt (Taylor-Entwicklung 2. Ordnung):

$$f(x+h) \approx f(x) + \langle \nabla f(x), h \rangle + \frac{1}{2} \langle h, \nabla^2 f(x) h \rangle$$

mit

$$\nabla^2 f(x) := \begin{pmatrix} \frac{\partial^2 f}{\partial x_1^2} & \frac{\partial^2 f}{\partial x_1 \partial x_2} & \cdots & \frac{\partial^2 f}{\partial x_1 \partial x_n} \\ \frac{\partial^2 f}{\partial x_2 \partial x_1} & \frac{\partial^2 f}{\partial x_2^2} & \cdots & \frac{\partial^2 f}{\partial x_2 \partial x_n} \\ \vdots & & & \vdots \\ \frac{\partial^2 f}{\partial x_n \partial x_1} & \cdots & \cdots & \frac{\partial^2 f}{\partial x_n^2} \end{pmatrix}$$

$\nabla^2 f(x)$ heißt auch „Hesse-Matrix“ (Diese ist immer symmetrisch!).

$$f(x+h) = f(x) + \underbrace{\sum_{i=1}^n \frac{\partial f}{\partial x_i}(x) \cdot h_i}_{\langle \nabla f(x), h \rangle} + \underbrace{\frac{1}{2} \sum_{i,j=1}^n h_i \frac{\partial^2 f_i}{\partial x_j}(x) h_j}_{=\langle h, \nabla^2 f(x) h \rangle}$$

(Analogie zur Taylor-Reihe 2. Ordnung für die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$!)

Beispiel 7.5.2. Beispiel für $f : \mathbb{R}^2 \rightarrow \mathbb{R}$

$$f(x) := (x_1 + 2x_2) + e^{-x_1} + e^{-x_2}$$

$$\nabla f(x) = \begin{pmatrix} \frac{\partial f}{\partial x_1} \\ \frac{\partial f}{\partial x_2} \end{pmatrix} = \begin{pmatrix} 1 - e^{-x_1} \\ 2 - e^{-x_2} \end{pmatrix}$$

$$\nabla^2 f(x) = \begin{pmatrix} \frac{\partial^2 f}{\partial x_1^2} & \frac{\partial^2 f}{\partial x_1 \partial x_2} \\ \frac{\partial^2 f}{\partial x_2 \partial x_1} & \frac{\partial^2 f}{\partial x_2^2} \end{pmatrix} = \begin{pmatrix} e^{-x_1} & 0 \\ 0 & e^{-x_2} \end{pmatrix}$$

Taylor-Entwicklung 2. Ordnung:

$$f(x+h) \approx f(x) + \langle \nabla f(x), h \rangle + \frac{1}{2} \langle h, \nabla^2 f(x) h \rangle$$

$$\begin{aligned}
&= f(x) + \begin{pmatrix} 1 - e^{-x_1} \\ 2 - e^{-x_2} \end{pmatrix} \cdot \begin{pmatrix} h_1 \\ h_2 \end{pmatrix} + \frac{1}{2} \underbrace{(h_1, h_2)}_{=h^T} \cdot \begin{pmatrix} e^{-x_1} & 0 \\ 0 & e^{-x_2} \end{pmatrix} \cdot \begin{pmatrix} h_1 \\ h_2 \end{pmatrix} \\
&= f(x) + (1 - e^{-x_1})h_1 + (2 - e^{-x_2})h_2 + \frac{1}{2}(h_1^2 e^{-x_1} + h_2^2 e^{-x_2})
\end{aligned}$$

Definition 7.5.3. Eine symmetrische Matrix $A \in \mathbb{R}^{n \times n}$ heißt

- (a) positiv definit, falls $\langle \xi, A\xi \rangle > 0 \quad \forall \xi \in \mathbb{R}^n, \xi \neq 0$
- (b) negativ definit, falls $\langle \xi, A\xi \rangle < 0 \quad \forall \xi \in \mathbb{R}^n, \xi \neq 0$
- (c) indefinit, falls es $\xi, \eta \in \mathbb{R}^n$ gibt mit $\langle \xi, A\xi \rangle > 0, \langle \eta, A\eta \rangle < 0$

Die lineare Algebra sagt uns:

- (a) A pos. definit $\Leftrightarrow$ alle Eigenwerte (EW) von A sind größer Null
- (b) A neg. definit $\Leftrightarrow$ alle EW sind kleiner Null
- (c) A indefinit $\Leftrightarrow A$ hat mindestens einen positiven und einen negativen EW

Wiederholung: Eigenwerte und Eigenvektoren

$$\begin{aligned}
|A - \lambda I| &= 0 & Av &= \lambda v \\
\downarrow & & & \\
Av - \lambda v &= 0 \\
\underbrace{(A - \lambda I)}_B v &= 0 \\
Bv &= 0
\end{aligned}$$

Plausibilitätserläuterung für Fall a):

Sei λ_1 positiver EW und $v^1 = \begin{pmatrix} v_1^1 \\ \vdots \\ v_n^1 \end{pmatrix} \in \mathbb{R}^n$ zugehöriger Eigenvektor (EV) von A , dann gilt:

$$\langle v^1, Av^1 \rangle = \langle v^1, \lambda_1 v^1 \rangle = \lambda_1 \langle v^1, v^1 \rangle = \lambda_1 \cdot \sum_{i=1}^n (v_i^1)^2 > 0, \text{ da } v^1 \neq 0 \quad (7.6)$$

Weitere Argumentation:

Da A symmetrisch $\exists$ Basis des $\mathbb{R}^n$ aus Eigenvektoren. Stelle beliebiges $\xi \in \mathbb{R}^n$ als Linearkombination dieser Basis dar und benutze 7.6, um Aussage a) zu verifizieren.

Richtungsableitung

Definition 7.5.4. Sei $U \subset \mathbb{R}^n$ offen, $f : U \rightarrow \mathbb{R}$ eine Funktion. Sei $u \in \mathbb{R}^n$ ein beliebiger Vektor der Länge (Norm) 1: $\|v\| = 1$. Die Richtungsableitung von f in einem Punkt $x \in U$ in Richtung v ist der Differentialquotient

$$\frac{d}{dt}f(x + tv)\Big|_{t=0} := \lim_{t \rightarrow 0} \frac{f(x + tv) - f(x)}{t},$$

$t \in \mathbb{R}$.

Satz 7.5.5. Sei $U \subset \mathbb{R}^n$ offen, $f : U \rightarrow \mathbb{R}$ stetig differenzierbar. Dann gilt $\forall x \in U$ und $\forall v \in \mathbb{R}^n$ mit $\|v\| = 1$:

$$\frac{d}{dt}f(x + tv)\Big|_{t=0} = \langle v, \nabla f(x) \rangle$$

Das Skalarprodukt $\langle v, \nabla f(x) \rangle$ beschreibt die Projektion des Gradienten ∇f von f auf die Richtung v :

$$\langle v, \nabla f(x) \rangle = \|v\| \cdot \|\nabla f(x)\| \cdot \cos \theta,$$

mit $\theta = \angle(v, \nabla f(x))$.

Beispiel Richtungsableitung

$$f(x, y) := xy^3 + x^2y^2, \quad f : \mathbb{R}^2 \rightarrow \mathbb{R}$$

Berechne Richtungsableitung von f im Punkt $(x_0, y_0) = (1, 1)$ in Richtung des Vektors $v = \begin{pmatrix} 4 \\ 3 \end{pmatrix}$:

1.

$$\nabla f(x, y) = \begin{pmatrix} \frac{\partial f}{\partial x}(x, y) \\ \frac{\partial f}{\partial y}(x, y) \end{pmatrix} = \begin{pmatrix} y^3 + 2xy^2 \\ 2x^2y + 3xy^2 \end{pmatrix}, \quad \nabla f(1, 1) = \begin{pmatrix} 3 \\ 5 \end{pmatrix}$$

2.

$$\|v\| = \sqrt{v_1^2 + v_2^2} = \sqrt{4^2 + 3^2} = 5$$

$\rightarrow w := \frac{v}{\|v\|}$ zeigt in Richtung v und hat die Länge (Norm) $\|w\| = 1$.

3. Berechne Richtungsableitung gemäß Satz 7.4.6:

$$\langle \nabla f(1, 1), w \rangle = \begin{pmatrix} 3 \\ 5 \end{pmatrix} \cdot \underbrace{\frac{1}{5} \begin{pmatrix} 4 \\ 3 \end{pmatrix}}_{= \frac{1}{\|v\|} \cdot v} = w = \frac{3 \cdot 4}{5} + \frac{5 \cdot 3}{5} = \frac{12}{5} + \frac{15}{5} = \frac{27}{5}$$

Satz 7.5.6. Sei $U \subset \mathbb{R}^n$ offen, $f : U \rightarrow \mathbb{R}$ differenzierbar. Falls f in $\bar{x} \in U$ ein lokales Extremum hat, so gilt die notwendige Bedingung

$$\nabla f(\bar{x}) = 0$$

Mann nennt $\bar{x} \in U$ dann auch „kritischen Punkt“ von f .

Das lokale Verhalten einer Funktion $f : U \rightarrow \mathbb{R}$ in der Umgebung eines kritischen Punktes wird durch ihre Hesse-Matrix $\nabla^2 f(\bar{x})$ bestimmt.

Erläuterung: Taylor-Entwicklung 2. Ordnung

$$f(\bar{x} + h) \approx f(\bar{x}) + \underbrace{\langle \nabla f(\bar{x}), h \rangle}_{=0, \text{ da } \nabla f(\bar{x})=0} + \frac{1}{2} \langle h, \nabla^2 f(\bar{x}) h \rangle$$

$$f(\bar{x}) + \frac{1}{2} \langle h, \nabla^2 f(\bar{x}) h \rangle,$$

Das macht folgende hinreichenden Bedingungen plausibel:

$\nabla^2 f(\bar{x})$ positiv definit $\Rightarrow \bar{x}$ lokales Minimum

$\nabla^2 f(\bar{x})$ negativ definit $\Rightarrow \bar{x}$ lokales Maximum

$\nabla^2 f(\bar{x})$ indefinit $\Rightarrow \bar{x}$ kein lokales Extremum

Die Eigenwertstruktur der Hesse-Matrix bestimmt also das lokale Extremalverhalten.

Beispiel 7.5.7. $f(x, y) = x_1^2 + x_2^2 + 2$

$$\nabla f(x_1, x_2) = \begin{pmatrix} \frac{\partial f}{\partial x_1}(x_1, x_2) \\ \frac{\partial f}{\partial x_2}(x_1, x_2) \end{pmatrix} = \begin{pmatrix} 2x_1 \\ 2x_2 \end{pmatrix}$$

$\rightarrow (x_1, x_2) = (0, 0)$ ist kritischer Punkt.

$$\nabla^2 f(x_1, x_2) = \begin{pmatrix} \frac{\partial^2 f}{\partial x_1^2} & \frac{\partial^2 f}{\partial x_1 \partial x_2} \\ \frac{\partial^2 f}{\partial x_2 \partial x_1} & \frac{\partial^2 f}{\partial x_2^2} \end{pmatrix} = \begin{pmatrix} 2 & 0 \\ 0 & 2 \end{pmatrix}$$

Eigenwerte der Hesse-Matrix:

$$\lambda_1 = 2 > 0, \lambda_2 = 2 > 0$$

$\Rightarrow \nabla^2 f(x_1, x_2)$ positiv definit $\rightarrow (x_1, x_2) = (0, 0)$ lokales Minimum !

im allgemeinen Fall: berechne EW von $\nabla^2 f$ über $|\nabla^2 f - \lambda I| = 0 \dots!$

Die Funktion aus dem letzten Beispiel beschreibt ein nach oben geöffnetes Paraboloid.

7.6 Funktionen vom $\mathbb{R}^n$ in den $\mathbb{R}^m$

Wir sind nun in der Lage, den allgemeinen Fall vektorwertiger Funktionen mit mehreren Argumenten zu betrachten, also Funktionen $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$, die z.B. zur Beschreibung von elektrischen Feldern, Strömungen, und dynamischen Systemen (siehe Kapitel 8) benötigt werden.

7.6.1 Allgemeiner Ableitungsbegriff im $\mathbb{R}^n$

Definition 7.6.1. Sei $U \subset \mathbb{R}^n$ offen und $f : U \rightarrow \mathbb{R}^m$ eine Abbildung. f heißt in $x \in U$ (total) differenzierbar, falls es eine lineare Abbildung $A : \mathbb{R}^n \rightarrow \mathbb{R}^m$ (siehe Kap. 7.1) gibt, so dass in einer Umgebung von x gilt:

$$f(x+h) = f(x) + A(h) + \varphi(h), \quad \text{mit } \lim_{h \rightarrow 0} \frac{\varphi(h)}{\|h\|} = 0$$

$$f : \mathbb{R}^n \rightarrow \mathbb{R}^m, (x_1, \dots, x_n) \mapsto \begin{pmatrix} f_1(x_1, \dots, x_n) \\ \vdots \\ f_m(x_1, \dots, x_n) \end{pmatrix}$$

mit den Komponentenfunktionen $f_i : \mathbb{R}^n \rightarrow \mathbb{R}$.

Die lineare Abbildung $A(h)$ kann bezüglich der Standardbasis aus Einheitsvektoren mit Hilfe einer Matrix beschrieben werden:

Definition 7.6.2 (Jacobi-Matrix).

Sei $U \subset \mathbb{R}^n$ offen, und $f : U \rightarrow \mathbb{R}^m$, und alle Komponenten $f_1, \dots, f_m$ stetig differenzierbar. Dann nennt man f stetig differenzierbar und schreibt $f \in C^1(U, \mathbb{R}^m)$. Die $(m \times n)$ -Matrix

$$f'(x) := \begin{pmatrix} \frac{\partial f_1}{\partial x_1}(x) & \frac{\partial f_1}{\partial x_2}(x) & \cdots & \frac{\partial f_1}{\partial x_n}(x) \\ \vdots & & & \vdots \\ \frac{\partial f_m}{\partial x_1}(x) & \cdots & \cdots & \frac{\partial f_m}{\partial x_n}(x) \end{pmatrix}$$

heißt die **Ableitung** oder die **Jacobi-Matrix** von f an der Stelle x . Manchmal schreibt man auch $\frac{\partial f}{\partial x}(x)$ statt $f'(x)$.

D.h. in Umgebung eines Punktes x gilt:

$$f(x+h) \approx f'(x) \cdot h, \quad h \in \mathbb{R}^n$$

mit $f'(x) \cdot h$ ein Matrix-Vektor-Produkt.

Ebenso wie für skalare Funktionen mehrerer Argumente gilt für vektorwertige Funktionen mehrerer Argumente eine Approximationseigenschaft analog zu Satz 7.4.10:

Satz 7.6.3 (Approximationseigenschaft der Jacobi-Matrix).

Sei $U \subset \mathbb{R}^n$ offen und $f \in C^1(U, \mathbb{R}^m)$. Dann gilt für alle $x, \bar{x} \in U$ dass

$$f(x) = f(\bar{x}) + f'(\bar{x})(x - \bar{x}) + \underbrace{\varphi(\bar{x}, x)}_{\text{Fehlervektor}}, \quad \text{wobei} \quad \lim_{x \rightarrow \bar{x}} \frac{\|\varphi(\bar{x}, x)\|}{\|x - \bar{x}\|} = 0.$$

Man beachte, dass $f'(\bar{x})$ eine Matrix ist und somit als lineare Abbildung $f'(x_0) : \mathbb{R}^n \rightarrow \mathbb{R}^m$ aufgefasst werden kann. In Abbildung 7.11 ist der Satz für den Fall $n = m = 1$ illustriert.

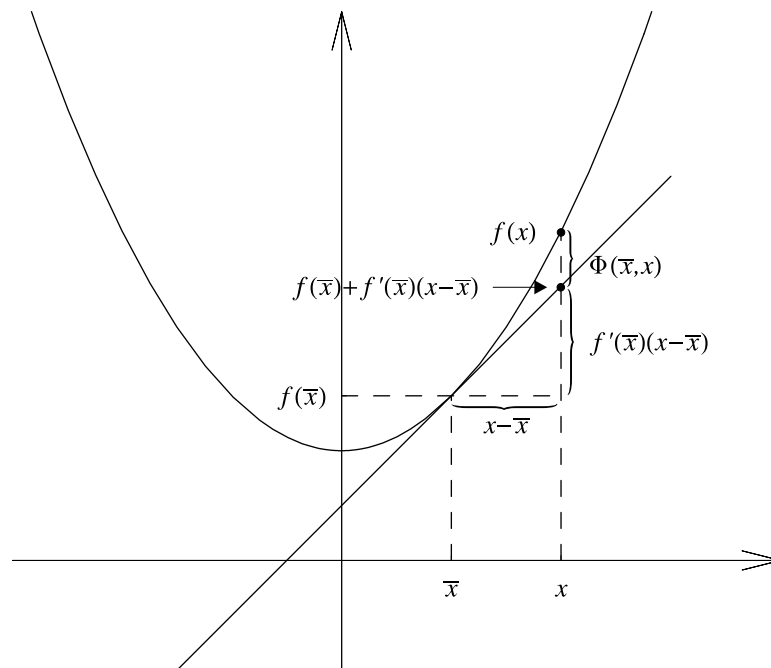


Abbildung 7.11: Approximationseigenschaft der Ableitung

Bemerkung 7.6.4.

Falls $n = m = 1$ ist besteht die „Jacobi-Matrix“ nur aus einer einzigen Zahl, und ist gerade die altbekannte Ableitung $f'(x)$ aus der Schule. Dies motiviert die Verwendung des Symbols $f'(x)$ zur Bezeichnung der Jacobi-Matrix. Ausserdem gilt

- Falls $m = 1$ ist die Jacobi-Matrix

$$f'(x) = \left(\frac{\partial f}{\partial x_1}(x), \dots, \frac{\partial f}{\partial x_n}(x) \right) = \left(\nabla f(x) \right)^T$$

der transponierte Gradient.

- Falls $n = 1$ ist die Jacobi-Matrix

$$f'(x) = \begin{pmatrix} \frac{\partial f_1}{\partial x}(x) \\ \vdots \\ \frac{\partial f_m}{\partial x}(x) \end{pmatrix}$$

der Tangentialvektor.

- Falls $n > 1$ und $m > 1$ besteht die Jacobi-Matrix $f'(x)$ aus übereinandergestapelten transponierten Gradienten der Einzelkomponentenfunktionen:

$$f(x) = \begin{pmatrix} f_1(x) \\ \vdots \\ f_n(x) \end{pmatrix} \Rightarrow f'(x) = \begin{pmatrix} f'_1(x) \\ \vdots \\ f'_n(x) \end{pmatrix} = \begin{pmatrix} \nabla f_1(x)^T \\ \vdots \\ \nabla f_n(x)^T \end{pmatrix}.$$

Beispiel 7.6.5 (Polarkoordinaten). Wir betrachten die Abbildung

$$\begin{aligned} f : \mathbb{R}^2 &\rightarrow \mathbb{R}^2 \\ x = \begin{pmatrix} r \\ \phi \end{pmatrix} &\mapsto f(x) = \begin{pmatrix} f_1(r, \phi) \\ f_2(r, \phi) \end{pmatrix} = \begin{pmatrix} r \sin \phi \\ r \cos \phi \end{pmatrix} \end{aligned}$$

und berechnen die partiellen Ableitungen

$$\begin{aligned} \frac{\partial f_1}{\partial r}(r, \phi) &= \sin \phi & \frac{\partial f_1}{\partial \phi}(r, \phi) &= r \cos \phi \\ \frac{\partial f_2}{\partial r}(r, \phi) &= \cos \phi & \frac{\partial f_2}{\partial \phi}(r, \phi) &= -r \sin \phi \end{aligned}$$

daraus ergibt sich die Jacobi-Matrix als

$$f'(x) = \frac{\partial f}{\partial (r, \phi)}(r, \phi) = \begin{pmatrix} \sin \phi & r \cos \phi \\ \cos \phi & -r \sin \phi \end{pmatrix}.$$

Die Verwendung von Jacobi-Matrizen wird besonders praktisch bei verknüpften Funktionen, denn es gilt eine verallgemeinerte Form der Kettenregel.

Satz 7.6.6 (Kettenregel für Jacobi-Matrizen).

Seien $f \in C^1(\mathbb{R}^n, \mathbb{R}^m)$ und $g \in C^1(\mathbb{R}^p, \mathbb{R}^n)$, dann ist auch ihre Verknüpfung stetig ableitbar, $f \circ g \in C^1(\mathbb{R}^p, \mathbb{R}^m)$, und es gilt

$$\underbrace{(f \circ g)'(x)}_{m \times p \text{ - Matrix}} = \underbrace{f'(g(x))}_{m \times n \text{ - Matrix}} \cdot \underbrace{g'(x)}_{n \times p \text{ - Matrix}}$$

Wir illustrieren den Satz anhand eines Beispiels.

Beispiel 7.6.7 (Logarithmische Spirale). Wir verknüpfen die Funktion f aus dem vorherigen Beispiel 7.6.5 mit einer Kurve

$$\begin{aligned} g : \mathbb{R} &\rightarrow \mathbb{R}^2 \\ t &\mapsto g(t) := \begin{pmatrix} e^t \\ t \end{pmatrix} \end{aligned}$$

Die Verknüpfung $(f \circ g)(t) = \begin{pmatrix} e^t \sin t \\ e^t \cos t \end{pmatrix}$ ist wieder eine Kurve, die in Abbildung 7.12 gezeigt ist. Mit $g'(t) = \begin{pmatrix} e^t \\ 1 \end{pmatrix}$ und Satz 7.6.6 gilt nun:

$$\begin{aligned} (f \circ g)'(t) &= f'(g(t)) \cdot g'(t) \\ &= \begin{pmatrix} \sin t & e^t \cos t \\ \cos t & -e^t \sin t \end{pmatrix} \cdot \begin{pmatrix} e^t \\ 1 \end{pmatrix} = \begin{pmatrix} e^t \sin t + e^t \cos t \\ e^t \cos t - e^t \sin t \end{pmatrix}. \end{aligned}$$

Zum Test leiten wir $f \circ g$ noch einmal direkt ab:

$$(f \circ g)(t) = \begin{pmatrix} e^t \sin t \\ e^t \cos t \end{pmatrix} \Rightarrow (f \circ g)'(t) = \begin{pmatrix} e^t \cos t + e^t \sin t \\ -e^t \sin t + e^t \cos t \end{pmatrix}.$$

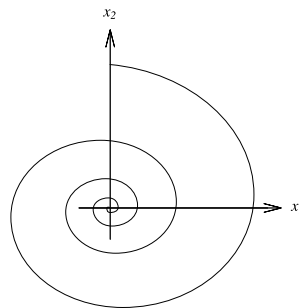


Abbildung 7.12: Die Kurve $f \circ g$ aus Beispiel 7.6.7.

7.7 Integration im $\mathbb{R}^n$

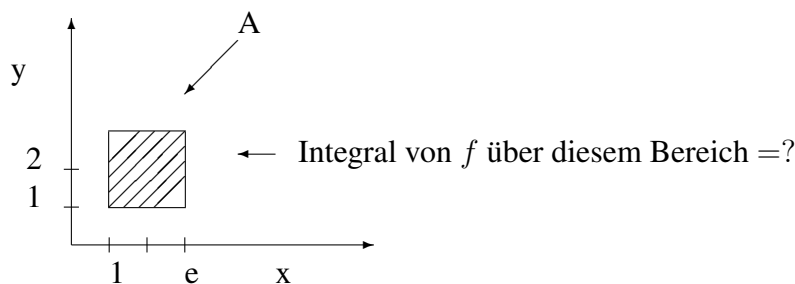
Es gibt hier zwei generelle Typen von Integralen:

1. Kurvenintegrale
2. Bereichsintegrale/Mehrfachintegrale

Typ 1 wird in Kapitel 7.9 (Vektoranalysis) behandelt. Typ 2 benutzt Maße für Mengen im $\mathbb{R}^n$, über denen reellwertige Funktionen integriert werden (zum Maßbegriff und Lebesgue-Integral siehe Kap. 6.2).

Beispiel 7.7.1 (Bereichsintegral).

$$f(x, y) = \frac{y^2}{x} \quad \int_A \int f(x, y) dA = ?$$



$$\begin{aligned} \int_1^2 \left(\int_1^e f(x, y) dx \right) dy &= \int_1^2 ([y^2 \ln x]_1^e) dy \\ &= \int_1^2 (y^2 \cdot 1 - 0) dy = \int_1^2 y^2 dy = \left[\frac{1}{3} y^3 \right]_1^2 = \frac{8}{3} - \frac{1}{3} = \frac{7}{3} \end{aligned}$$

Häufig wollen wir Integrale von Funktionen berechnen, die nicht nur von einem einzigen Argument abhängen, sondern von mehreren $x_1, \dots, x_n$. Als Beispiel sei z.B. die Gesamtmasse eines Körpers K mit ortsabhängiger Dichte $\rho(x) = \rho(x_1, x_2, x_3)$ genannt. Die Integration soll dann nicht auf einem Integrationsintervall, wie im Falle $n = 1$ stattfinden, sondern auf einem „Integrationsvolumen“, im Beispiel wäre dies z.B. $K \subset \mathbb{R}^3$. Man könnte also die Gesamtmasse m durch einen Ausdruck der Form

$$m := \int_K \rho(x) dV$$

beschreiben, wobei K das Volumen des Körpers und dV ein infinitesimales Volumenelement ist. Es sollen sozusagen die infinitesimalen Massenstücke $\rho(x)dV$ über alle Orte $x \in K$ aufsummiert werden. Es stellen sich im Wesentlichen zwei Probleme für die konkrete Berechnung solcher Integrale:

- Wie beschreibt man das Integrationsgebiet?
- Wie beschreibt man das infinitesimale Volumenelement?

Diese beiden Fragen werden wir im folgenden Abschnitt für einige für Sie wichtige Spezialfälle beantworten.

7.7.1 Sukzessive Integration

Am einfachsten ist der Fall eines quaderförmigen Integrationsgebiet im $\mathbb{R}^n$, also einer Menge $M = I_1 \times \cdots \times I_n$, wobei I_k Intervalle sind. Das Prinzip wird schon im Falle $n = 2$ deutlich, wo also eine Funktion $f(x_1, x_2)$ über eine Fläche $I_1 \times I_2 = [a, b] \times [c, d]$ integriert werden soll. Es gilt der folgende Satz.

Satz 7.7.2 (Sukzessive Integration).

Sei $f : [a, b] \times [c, d] \rightarrow \mathbb{R}$ stetig. Dann existieren die beiden Doppel-Integrale

$$\int_a^b \left(\int_c^d f(x_1, x_2) dx_2 \right) dx_1 = \int_c^d \left(\int_a^b f(x_1, x_2) dx_1 \right) dx_2,$$

und sind einander gleich, d.h. die Integrationsreihenfolge ist egal.

Man schreibt auch

$$\int_a^b \int_c^d f(x_1, x_2) \underbrace{d^2x}_{\text{2-D-Element}} \quad \text{oder} \quad \int_{[a,b] \times [c,d]} f(x_1, x_2) d^2x.$$

Beispiel 7.7.3.

$$\int_2^3 \left(\int_0^1 \cos(x_1 x_2) dx_2 \right) dx_1 = \int_2^3 \frac{1}{x_1} \sin(x_1 x_2) \Big|_0^1 dx_1 = \int_2^3 \left(\frac{1}{x_1} \sin(x_1) \right) dx_1.$$

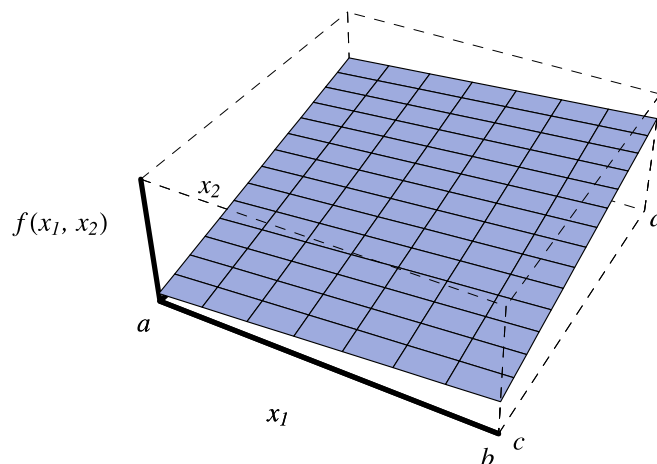


Abbildung 7.13: Das Integral als Volumen unter dem Graphen.

Bemerkung 7.7.4. Das Integral über einer Fläche kann als *Volumen* unter dem Graphen interpretiert werden, siehe Abbildung 7.13.

Integration auf einem gekrümmten Gebiet

Integration über einem Rechteck ist also einfach. Aber was ist, wenn statt über einem Rechteck $[a, b] \times [c, d]$ über eine Menge der Form

$$M = \{x \in \mathbb{R}^2 \mid x_1 \in [a, b], g_1(x_1) \leq x_2 \leq g_2(x_1)\}$$

mit stetigen Funktionen g_1, g_2 integriert werden soll (siehe Abbildung 7.14)? Dafür zerschneidet man die Fläche in senkrechte Streifen und berechnet wieder sukzessive ein Doppel-Integral

$$\int_a^b \left(\int_{g_1(x_1)}^{g_2(x_1)} f(x_1, x_2) dx_2 \right) dx_1.$$

Achtung: hier kann die Reihenfolge *nicht* vertauscht werden, da da x_1 in der Grenze des inneren Integrals vorkommt!

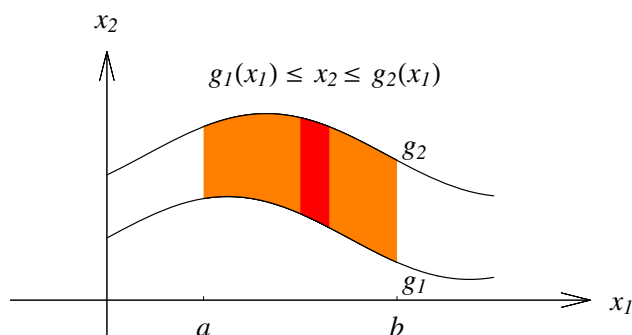


Abbildung 7.14: Integrationsfläche zwischen zwei Funktionen g_1 und g_2 .

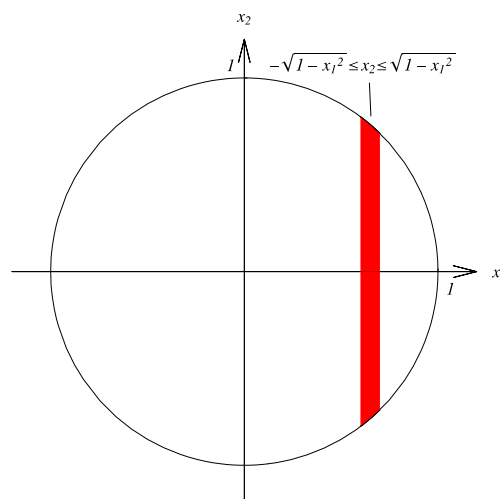


Abbildung 7.15: Zerlegung des Kreises in senkrechte Streifen

Beispiel 7.7.5 (Integration auf einer Kreisscheibe). Wir möchten

$$\int_M f(x_1, x_2) dx_1 dx_2 \quad \text{berechnen}$$

$$\text{wobei } M := \{x \in \mathbb{R}^2 \mid x_1^2 + x_2^2 \leq 1\}.$$

Hier hilft am Aufmalen der Menge M wie in Abbildung 7.15 und „Zerschneiden“ in senkrechte Streifen. Das Ergebnis ist ein Doppel-Integral, von dem wir im Prinzip wissen, wie wir es ausrechnen können:

$$\int_M f(x_1, x_2) \, dx_1 dx_2 = \int_{-1}^1 \left(\int_{-\sqrt{1-x_1^2}}^{\sqrt{1-x_1^2}} f(x_1, x_2) \, dx_2 \right) dx_1.$$

7.8 Integration in verschiedenen Koordinatensystemen

Oft ist es praktisch, statt in rechtwinkligen, kartesischen Koordinaten in einem anderen Koordinatensystem zu integrieren. Dies ist z.B. der Fall, wenn man weiss, dass die zu integrierende Funktion nur vom Abstand zum Ursprung abhängt, oder wenn das Integrationsgebiet z.B. kugelförmig ist. Wir beginnen hier mit einem wichtigen Spezialfall zur Motivation, der Integration in Polarkoordinaten, und geben dann eine allgemeine Regel für die Integration nach Koordinatentransformationen an, die wir noch auf einen weiteren Spezialfall anwenden.

7.8.1 Polarkoordinaten

In Polarkoordinaten stellen wir einen Vektor $x \in \mathbb{R}^2$ durch seinen Abstand zum Ursprung (oder Radius) r und durch den Winkel ϕ dar, den er mit der x_1 -Achse im mathematisch positiven Sinne bildet. Die Vorschrift, die jedem Paar (r, ϕ) einen Vektor (x_1, x_2) zuordnet, ist durch die bijektive Funktion

$$g :]0, \infty[\times]-\pi, \pi[\rightarrow \mathbb{R}^2 \setminus \{0\} \quad (7.7)$$

$$\begin{pmatrix} r \\ \phi \end{pmatrix} \mapsto g(r, \phi) = \begin{pmatrix} r \cos \phi \\ r \sin \phi \end{pmatrix}$$

gegeben, wobei wir den Nullvektor im Ursprung ($r = 0$) weglassen haben, da wir ihm keinen Winkel ϕ zuordnen könnten und sonst die Bijektivität aufgeben müssten.

Das Schöne an einer solchen bijektiven Koordinatentransformation ist nun, dass jede Funktion $f(x)$ von $x = (x_1, x_2)$ auch als Funktion $\tilde{f}(r, \phi)$ von (r, ϕ) dargestellt werden kann, nämlich durch $\tilde{f}(r, \phi) := f(g(r, \phi))$. Umgekehrt gilt natürlich auch $f(x) = \tilde{f}(g^{-1}(x))$.

Beispiel 7.8.1 (Höhenprofil eines Ameisenhaufens).

Wir betrachten das Höhenprofil eines Ameisenhaufens, siehe Abbildung 7.16, dessen Mittelpunkt im Ursprung liegt. Die jeweilige Höhe $h(x_1, x_2)$ sei als Funktion vom Ort (x_1, x_2) auf der Grundfläche wie folgt gegeben:

$$h(x_1, x_2) := H - H \frac{x_1^2 + x_2^2}{R^2} \quad \text{mit} \quad x_1^2 + x_2^2 \leq R^2.$$

In Polarkoordinaten ergibt sich der einfachere Ausdruck:

$$\tilde{h}(r, \phi) = h(g(r, \phi)) = H - H \frac{(r \cos \phi)^2 + (r \sin \phi)^2}{R^2} = H - H \frac{r^2}{R^2} \quad \text{mit} \quad r \leq R.$$

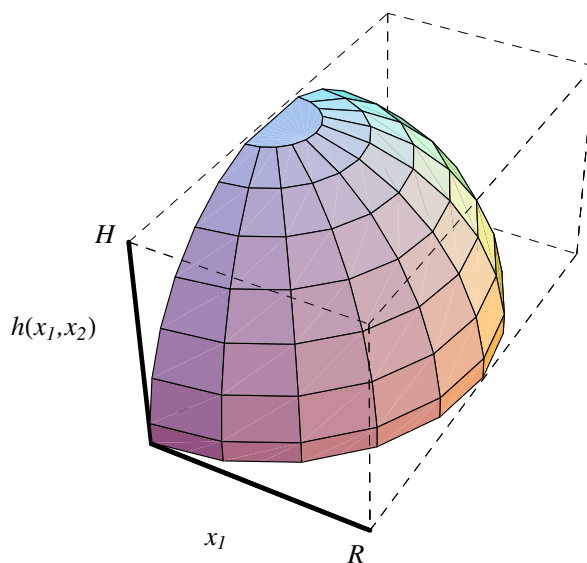


Abbildung 7.16: Höhenprofil des Ameisenhaufens

7.8.2 Integration in Polarkoordinaten

Um jetzt ein Integral

$$\int_M f(x) \, d^2x,$$

das in kartesischen Koordinaten gegeben ist, in Polarkoordinaten integrieren zu können, also unter Verwendung der Funktion $\tilde{f}(r, \phi) = f(g(r, \phi))$, müssen wir noch beantworten, wie wir

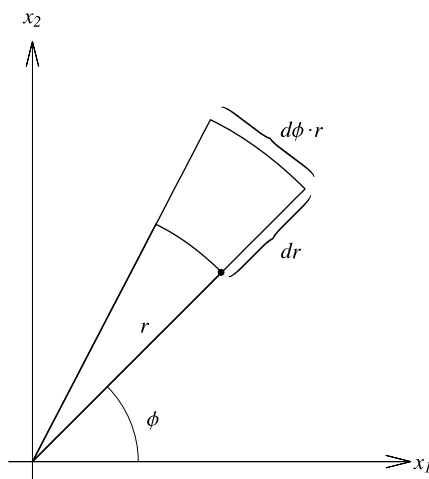
- das Integrationsgebiet M , und
- das infinitesimale Flächenstück d^2x

in Polarkoordinaten ausdrücken? Das neue Integrationsgebiet ist recht einfach als Urbild $g^{-1}(M)$ zu ermitteln, wie wir gleich am Beispiel illustrieren werden. Aber wie groß wird das Flächenstück d^2x , wenn wir im (r, ϕ) -Raum integrieren? Man kann sich durch geometrische Überlegungen davon überzeugen (siehe Abbildung 7.17), dass

$$d^2x = r \cdot d\phi \cdot dr. \quad (7.8)$$

Eine algebraische Herleitung dieser Identität, die aus einem allgemeinen Satz über Integration nach Koordinatentransformationen folgt, wird im folgenden Abschnitt gegeben. Dies erlaubt uns nun, das Integral in Polarkoordinaten auszudrücken als:

$$\int_M f(x) \, d^2x = \int_{g^{-1}(M)} \tilde{f}(r, \phi) \, r \, d\phi \, dr.$$

Abbildung 7.17: Das infinitesimale Flächenstück $r \cdot d\phi \cdot dr$ in Polarkoordinaten.**Beispiel 7.8.2** (Volumen eines Ameisenhaufens).

Um das Volumen des Ameisenhaufens aus Beispiel 7.8.1 zu berechnen, integrieren wir sein Höhenprofil $h(x) = H - H \frac{x_1^2 + x_2^2}{R^2}$ über seine Grundfläche, $F := \{x | x_1^2 + x_2^2 \leq R^2\}$, siehe Abbildung 7.16. Wir wollen also

$$\int_F h(x) d^2x = \int_{x_1^2 + x_2^2 \leq R^2} \left(H - H \frac{x_1^2 + x_2^2}{R^2} \right) d^2x \quad (7.9)$$

berechnen. Das Integrationsgebiet F transformieren wir zu

$$g^{-1}(F) = \{(r, \phi) | g(r, \phi) \in F\} = \{(r, \phi) | (r \cos \phi)^2 + (r \sin \phi)^2 \leq R^2\} = \{(r, \phi) | 0 < r \leq R\}.$$

So können wir jetzt das Integral (7.9) über einer Kreisscheibe in Polarkoordinaten viel einfacher berechnen als Integral über dem Quadrat $]0, R] \times [-\pi, \pi[$:

$$\begin{aligned} \int_0^R \int_{-\pi}^{\pi} \tilde{h}(r, \phi) r d\phi dr &= 2\pi \int_0^R \left(H - H \frac{r^2}{R^2} \right) r dr = 2\pi \int_0^R \left(Hr - \frac{Hr^3}{R^2} \right) dr \\ &= 2\pi \left(\frac{Hr^2}{2} - \frac{Hr^4}{4R^2} \right) \Big|_0^R = 2\pi \left(\frac{HR^2}{2} - \frac{HR^4}{4R^2} \right) \\ &= \frac{\pi}{2} HR^2. \end{aligned}$$

Beispiel 7.8.3 (Volumen der Gauss-Glocke).

Als ein zweites Beispiel für die Integration in Polarkoordinaten wollen wir die zweidimensionale Gauss-Glocke $f(x) = e^{-x_1^2 - x_2^2}$ über den gesamten $\mathbb{R}^2$ integrieren (vgl. Beispiel 9.2.9). Nach Transformation in Polarkoordinaten ergibt sich $\tilde{f}(r, \varphi) = e^{-r^2}$ und somit

$$\begin{aligned}
\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-x_1^2 - x_2^2} dx_1 dx_2 &= \int_0^{\infty} \int_{-\pi}^{\pi} e^{-r^2} \cdot r \cdot d\phi \cdot dr \\
&= \int_0^{\infty} 2\pi e^{-r^2} \cdot r \cdot dr \\
&= \int_0^{\infty} \pi e^{-y} dy \quad \left[y = r^2 \quad dy = 2r dr \right] \\
&= -\pi e^{-y} \Big|_0^{\infty} = 0 - (-\pi e^{-0}) = \pi.
\end{aligned}$$

Ganz nebenbei haben wir damit auch gleich das ansonsten sehr schwer zu berechnende Integral der eindimensionalen Gauss-Glocke ausgerechnet, $\int_{-\infty}^{\infty} e^{-x^2} dx = \sqrt{\pi}$, denn es gilt

$$\begin{aligned}
\left(\int_{-\infty}^{\infty} e^{-x^2} dx \right)^2 &= \left(\int_{-\infty}^{\infty} e^{-x_1^2} dx_1 \right) \cdot \left(\int_{-\infty}^{\infty} e^{-x_2^2} dx_2 \right) \\
&= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-x_1^2} \cdot e^{-x_2^2} dx_1 dx_2 \\
&= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-x_1^2 - x_2^2} dx_1 dx_2 = \pi.
\end{aligned}$$

7.9 *Integration nach Koordinatentransformationen

Ganz allgemein gilt bei der Integration einer Funktion f auf einem Gebiet $U \subset \mathbb{R}^n$ nach einer bijektiven Koordinatentransformation $g : W \rightarrow U$ folgender Satz:

Satz 7.9.1 (Integration nach Koordinatentransformationen).

Seien $U, W \subset \mathbb{R}^n$ offen, und $f \in C^0(U, \mathbb{R})$, sowie g eine bijektive Koordinatentransformation zwischen U und $W = g^{-1}(U)$ mit

$$g \in C^1(W, U) \text{ und } g^{-1} \in C^1(U, W).$$

Dann gilt

$$\int_U f(x) d^n x = \int_{W=g^{-1}(U)} f(g(y)) \underbrace{\det(g'(y))}_{\text{Determinante der Jacobi-Matrix}} d^n y.$$

Beispiel 7.9.2 (Flächenelement in Polarkoordinaten).

Zur Illustration des Satzes betrachten wir noch einmal die Integration in Polarkoordinaten. Hier gilt $y = (r, \phi)$, $U = \mathbb{R}^2 \setminus \{0\}$ und $W =]0, \infty[\times]-\pi, \pi[$. Die Funktion g aus (7.7) ist bijektiv, mit Umkehrabbildung

$$g^{-1} : U \rightarrow W \quad (7.10)$$

$$x \mapsto g^{-1}(x) = \begin{pmatrix} \sqrt{x_1^2 + x_2^2} \\ \arcsin \frac{x_2}{\sqrt{x_1^2 + x_2^2}} \end{pmatrix} = \begin{pmatrix} r \\ \phi \end{pmatrix}.$$

Die Funktion g^{-1} ist sicher stetig ableitbar, jedoch sind wir nur an der Jacobi-Matrix von g selbst interessiert, um die Determinante $\det(g'(y))$ zu bestimmen:

$$\begin{aligned} g(r, \phi) = \begin{pmatrix} r \cos \phi \\ r \sin \phi \end{pmatrix} &\Rightarrow g'(r, \phi) = \begin{pmatrix} \cos \phi & -r \sin \phi \\ \sin \phi & r \cos \phi \end{pmatrix} \\ &\Rightarrow \det(g'(x)) = r \cos^2 \phi - -r \sin^2 \phi \\ &= r. \end{aligned}$$

Daraus ergibt sich eine nachträgliche algebraische Begründung für Behauptung (7.8):

$$d^2x = \det(g'(r, \phi)) \cdot d\phi \cdot dr = r \cdot d\phi \cdot dr.$$

7.9.1 *Integration in Kugelkoordinaten

Als eine weitere wichtige Koordinatentransformation wollen wir die Transformation des $\mathbb{R}^3$ in Kugelkoordinaten (oder auch „sphärische Koordinaten“) betrachten, die in Abbildung 7.18 illustriert ist. Es ist hilfreich, sich die Kugelkoordinaten mit Hilfe eines Globus mit Längen- und Breitengraden vorzustellen. Hier wird jeder Vektor $x \in \mathbb{R}^3$ durch den Ausdruck

$$\begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} r \sin \theta \cdot \cos \phi \\ r \sin \theta \cdot \sin \phi \\ r \cos \theta \end{pmatrix}$$

dargestellt, wobei

- $r \in]0, \infty[$ den Abstand vom Ursprung (den *Radius*) darstellt, sowie
- $\phi \in]-\pi, \pi[$ als *Längengrad* und
- $\theta \in]0, \pi[$ als *Breitengrad* bezeichnet werden kann.

In der Sprache des vorherigen Abschnitts definieren wir uns also eine bijektive Koordinatentransformation

$$g :]0, \infty[\times]-\pi, \pi[\times]0, \pi[\rightarrow \mathbb{R}^3 \setminus \{0\} \quad (7.11)$$

$$y = \begin{pmatrix} r \\ \phi \\ \theta \end{pmatrix} \mapsto g(y) := \begin{pmatrix} r \sin \theta \cdot \cos \phi \\ r \sin \theta \sin \phi \\ r \cos \theta \end{pmatrix}.$$

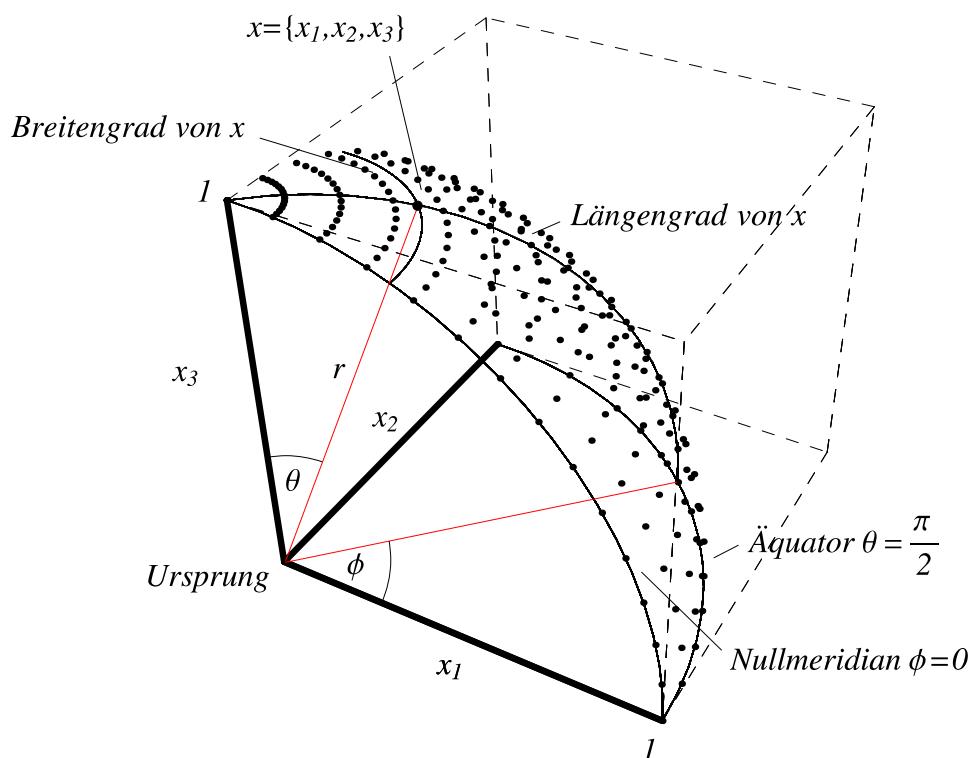


Abbildung 7.18: Kugelkoordinaten

Die etwas aufwendige Berechnung der Jacobi-Matrix $g'(y)$ und ihrer Determinante ergibt $\det(g'(y)) = r^2 \sin \theta$ und somit gilt für das dreidimensionale Volumenelement (analog zu (7.8)):

$$d^3x = r^2 \sin \theta \cdot d\phi \cdot r d\theta \cdot dr.$$

Demnach gilt mit Satz 7.9.1 für das Integral einer beliebigen Funktion f über einer Menge $M \subset \mathbb{R}^3$:

$$\int_M f(x) d^3x = \int_{g^{-1}(M)} f(g(r, \phi, \theta)) \cdot r^2 \sin \theta d\phi d\theta dr.$$

Beispiel 7.9.3 (Masse der Erdatmosphäre).

Welche Masse hat die Erdatmosphäre? Die Dichte $\rho(h)$ hängt nur von der Höhe über dem Erdboden ab, und R ist der Erdradius. Wir wollen über das gesamte Volumen oberhalb der Erdoberfläche integrieren, also über die Menge $M = \{x \in \mathbb{R}^3 \mid \|x\| > R\}$. Da die Höhe über dem Erdboden durch $\|x\| - R = \sqrt{x_1^2 + x_2^2 + x_3^2} - R$ gegeben ist, müssten wir in kartesischen Koordinaten das Integral

$$\int_{\left\{x \in \mathbb{R}^3 \mid \sqrt{x_1^2 + x_2^2 + x_3^2} > R\right\}} \rho\left(\sqrt{x_1^2 + x_2^2 + x_3^2} - R\right) d^3x$$

berechnen. In Kugelkoordinaten erhalten wir den wesentlich einfacheren Ausdruck

$$\begin{aligned} \int_R^\infty \int_{-\pi}^\pi \int_0^\pi \rho(r-R) \cdot r^2 \sin \theta \cdot d\theta \cdot d\phi \cdot dr &= \int_0^\infty \int_{-\pi}^\pi \rho(r-R) 2r^2 d\phi dr \\ \left[\text{denn } \int_0^\pi \sin \theta \cdot d\theta = -\cos \theta \Big|_0^\pi = 2 \right] & \\ &= \int_R^\infty \rho(r-R) \cdot 4\pi r^2 dr. \end{aligned}$$

Der Ausdruck $4\pi r^2 dr$ kann dabei als das Volumen einer infinitesimalen Kugelschale (eines „Zwiebelrings“) mit Dicke dr und Oberfläche $4\pi r^2$ interpretiert werden, siehe Abbildung 7.19.

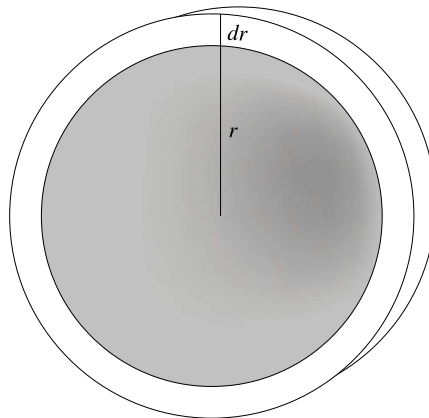


Abbildung 7.19: Eine (aufgeschnittene) infinitesimale Kugelschale.

Legen wir der Berechnung der Dichte ρ die barometrische Höhenformel $\rho(h) = \rho_0 \cdot e^{-\beta h}$ zugrunde, ergibt sich das Integral

$$4\pi\rho_0 \int_R^\infty r^2 e^{-\beta(r-R)} dr = e^{\beta R} \cdot 4\pi\rho_0 \int_R^\infty r^2 e^{-\beta r} dr = 4\pi\rho_0 \left(\frac{R^2}{\beta} + \frac{2R}{\beta^2} + \frac{2}{\beta^3} \right),$$

denn durch zweimaliges partielles Integrieren erhalten wir

$$\begin{aligned}
 \int_R^\infty r^2 e^{-\beta r} dr &= -\frac{1}{\beta} r^2 e^{-\beta r} \Big|_R^\infty + \int_R^\infty \frac{2r}{\beta} e^{-\beta r} \\
 &= -\frac{1}{\beta} r^2 e^{-\beta r} \Big|_R^\infty - \frac{2r}{\beta^2} e^{-\beta r} \Big|_R^\infty + \int_R^\infty \frac{2}{\beta^2} e^{-\beta r} \\
 &= -\frac{r^2}{\beta} e^{-\beta r} \Big|_R^\infty - \frac{2r}{\beta^2} e^{-\beta r} \Big|_R^\infty - \frac{2}{\beta^3} e^{-\beta r} \Big|_R^\infty \\
 &= e^{-\beta R} \left(\frac{R^2}{\beta} + \frac{2R}{\beta^2} + \frac{2}{\beta^3} \right).
 \end{aligned}$$

7.10 Kurskurs Optimierung im $\mathbb{R}^n$

Wir wollen uns in diesem sehr knappen Abschnitt kurz der Frage zuwenden, wie man eine Minimalstelle $x^* \in \mathbb{R}^n$ einer Funktion $f \in C^2(\mathbb{R}^n, \mathbb{R})$ finden und charakterisieren kann. Dafür rufen wir uns zunächst in Erinnerung, was wir schon über Minimalstellen von Funktionen eines Argumentes wissen: Dafür, dass $x^* \in \mathbb{R}$ lokale Minimalstelle einer Funktion $f \in C^2(\mathbb{R}, \mathbb{R})$ ist, gab es zwei Bedingungen:

1. Notwendige Bedingung: $f'(x^*) = 0$.
2. Hinreichende Bedingung: $f'(x^*) = 0$ und $f''(x^*) > 0$.

Dafür, dass $x^* \in \mathbb{R}^n$ lokale Minimalstelle einer Funktion $f \in C^2(\mathbb{R}^n, \mathbb{R})$ mit mehreren Argumenten ist, gelten nun ähnliche Bedingungen, die etwas komplexer sind:

1. Notwendige Bedingung: $\nabla f(x^*) = 0$.
2. Hinreichende Bedingung: $\nabla f(x^*) = 0$ und die sogenannte **Hesse-Matrix** $\nabla^2 f(x) = \left(\frac{\partial^2 f}{\partial x_i \partial x_j} \right) \in \mathbb{R}^{n \times n}$ hat nur positive Eigenwerte.

Bemerkung 7.10.1 (Hesse-Matrix als Jacobi-Matrix).

Es ist vielleicht hilfreich, sich klarzumachen, dass die Hesse-Matrix $\nabla^2 f(x)$ nichts anderes ist als die quadratische Jacobi-Matrix des Gradienten ∇f , der ja eine Funktion $\nabla f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ ist:

$$\nabla^2 f(x) = \left(\nabla f \right)'(x).$$

Da der Gradient aber selbst bereits aus ersten Ableitungen besteht, ist seine Jacobi-Matrix nach dem Satz 7.4.20 über die Vertauschbarkeit zweiter partieller Ableitungen symmetrisch.

Beispiel 7.10.2 (Minimalstelle im $\mathbb{R}^2$).

Wir suchen ein Minimum von der Funktion $f : \mathbb{R}^2 \rightarrow \mathbb{R}$, $f(x) := (x_1 + 2x_2) + e^{-x_1} + e^{-x_2}$.

1. Notwendige Bedingung: wir berechnen

$$\nabla f(x) = \begin{pmatrix} 1 - e^{-x_1} \\ 2 - e^{-x_2} \end{pmatrix}. \quad \text{Also ist} \quad \nabla f(x^*) = 0 \Leftrightarrow \begin{matrix} x_1^* = -\ln 1 \\ x_2^* = -\ln 2. \end{matrix}$$

2. Hinreichende Bedingung: für $x^* = \begin{pmatrix} -\ln 1 \\ -\ln 2 \end{pmatrix}$ müssen wir nur noch die Eigenschaft der Hesse-Matrix testen. Wir berechnen

$$\nabla^2 f(x) = \begin{pmatrix} \frac{\partial^2 f}{\partial x_1^2} & \frac{\partial^2 f}{\partial x_1 \partial x_2} \\ \frac{\partial^2 f}{\partial x_1 \partial x_2} & \frac{\partial^2 f}{\partial x_2^2} \end{pmatrix} = \begin{pmatrix} e^{-x_1} & 0 \\ 0 & e^{-x_2} \end{pmatrix}$$

An der Stelle x^* gilt also $\nabla^2 f(x^*) = \begin{pmatrix} 1 & 0 \\ 0 & 2 \end{pmatrix}$ mit positiven Eigenwerten 1 und 2.

Die notwendige Bedingung hilft uns beim Auffinden einer möglichen Minimalstelle x^* , und mit der hinreichenden Bedingung können wir prüfen, ob x^* tatsächlich Minimalstelle von f ist.

7.11 Vektoranalysis

Eine Abbildung $f : \mathbb{R}^n \supset U \rightarrow \mathbb{R}^n$ nennt man oft auch **Vektorfeld**. Im Falle $n = 2, 3$ schreibt man oft $\vec{x} \mapsto \vec{f}(\vec{x})$. Vektorfelder zeichnen sich dadurch aus, dass die Vektoren $f(x) \in \mathbb{R}^n$ selbst Elemente des gleichen Vektorraumes sind, in dem die Argumente $x \in \mathbb{R}^n$ liegen. Das ermöglicht z.B. die Beschreibung von Geschwindigkeitsfeldern oder Kraftfeldern; wir werden im Kapitel 8 über dynamische Systeme noch eine weitere äusserst wichtige Anwendung von Vektorfeldern kennenlernen.

Definition 7.11.1 (Zeitabhängiges und -unabhängiges Vektorfeld).

Sei $U \subset \mathbb{R}^n$ und $I \subset \mathbb{R}$ ein Intervall.

Eine Funktion $f : U \times I \rightarrow \mathbb{R}^n, (t, x) \mapsto f(t, x)$ nennt man ein **zeitabhängiges Vektorfeld**, und eine Funktion $f : U \rightarrow \mathbb{R}^n, x \mapsto f(x)$ nennt man **zeitunabhängiges Vektorfeld in $\mathbb{R}^n$** .

Beispiel 7.11.2. Wir geben hier einige Beispiele für Vektorfelder.

- Die Geschwindigkeit einer Wasserströmung in einem Rohr (siehe Abbildung 7.20), ist ein Vektorfeld im $\mathbb{R}^3$, denn an jedem Punkt $\vec{x} \in U \subset \mathbb{R}^3$ hat das Wasser eine Geschwindigkeit $\vec{v}(\vec{x}) \in \mathbb{R}^3$. Falls die Strömung **stationär** ist, wie z.B. die Strömung unter einem schwach aufgedrehten Wasserhahn, dessen Strahl stillsteht wie ein Eiszapfen, ist das Vektorfeld zeitunabhängig. Im Fall turbulenter Strömung mit wandernden Wirbeln ist das Vektorfeld zeitabhängig, und man müßte $\vec{v}(t, \vec{x})$ schreiben.
- Der Gradient ∇f einer skalaren Funktion $f \in C^1(\mathbb{R}^n, \mathbb{R})$ ist ein Vektorfeld im $\mathbb{R}^n$. Für die Funktion $f(x) = x_1^2 + x_2^2$ (siehe Abbildung 7.21) ist das Gradientenfeld

$$\nabla f : \mathbb{R}^2 \rightarrow \mathbb{R}^2, \quad x \mapsto \nabla f(x) = \begin{pmatrix} 2x_1 \\ 2x_2 \end{pmatrix}$$

beispielsweise ein zeitunabhängiges Vektorfeld im $\mathbb{R}^2$.

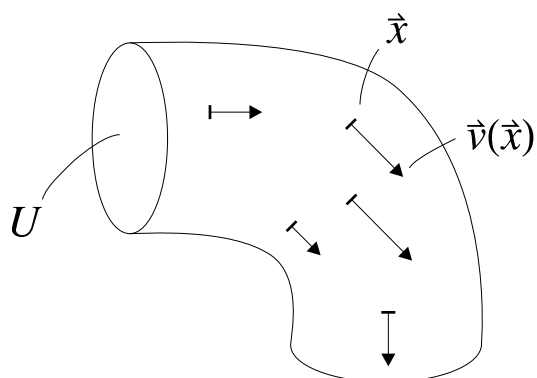


Abbildung 7.20: Stationäre Wasserströmung im Rohr

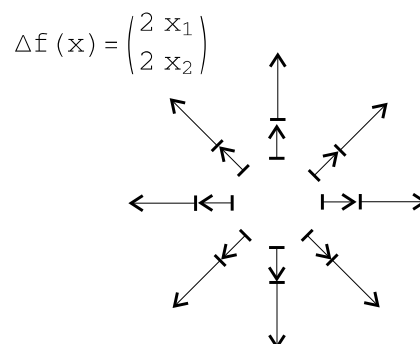
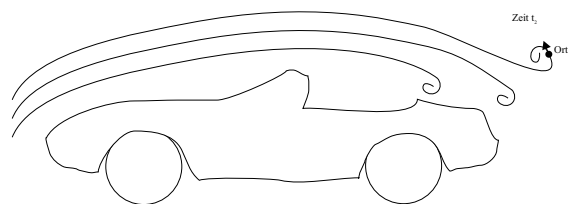
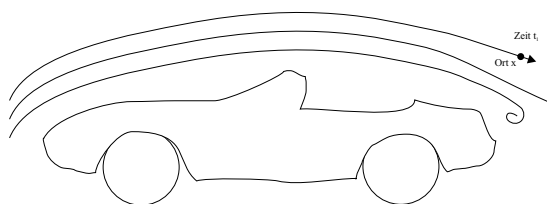
Abbildung 7.21: Das Gradientenfeld zu $f(x) = x_1^2 + x_2^2$ 

Abbildung 7.22: Auto im Windkanal, zu zwei verschiedenen Zeitpunkten.

- Windkanal (Abbildung 7.22): Auch hier hat die Luft zu jedem Zeitpunkt und an jedem Ort eine bestimmte Geschwindigkeit, und alle Geschwindigkeiten zusammen bilden wieder ein Vektorfeld im $\mathbb{R}^3$. Dieses Feld ist wegen der Turbulenz typischerweise zeitabhängig.
- Das elektrische Feld $\vec{E}(\vec{x})$ ist ein Vektorfeld im $\mathbb{R}^3$, siehe Abbildung 7.23. Falls die es erzeugenden Ladungen stillstehen, ist es zeitunabhängig. Die „Stromlinien“ nennt man hier **Feldlinien**. Man kann zeigen, dass ein stationäres elektrisches Feld das Gradientenfeld einer skalaren Funktion $\phi(\vec{x})$ darstellt, d.h. $\vec{E}(\vec{x}) = \nabla\phi(\vec{x})$. Diese Funktion ϕ nennt man das *elektrische Potential*. Auf Potentiale werden wir später noch etwas genauer eingehen.
- Das Gravitationsfeld der Erde ist ein Vektorfeld im $\mathbb{R}^3$, das wir als stationär ansehen können, wenn sich unser Koordinatensystem mit der Erde durch den Raum bewegt. Auch hier gibt es eine skalare Funktion, als dessen Gradient das Gravitationsfeld angesehen werden kann, das sogenannte *Gravitationspotential*.

Veranschaulichung von Vektorfeldern

Es gibt im wesentlichen zwei Möglichkeiten, sich ein Vektorfeld zu veranschaulichen. Beide Methoden funktionieren am besten im $\mathbb{R}^2$:

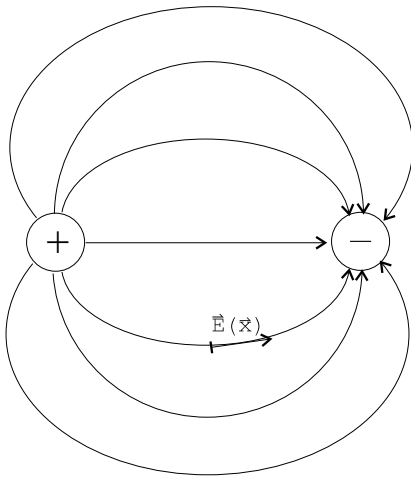


Abbildung 7.23: Das elektrische Feld zwischen zwei entgegengesetzt geladenen Kugeln.

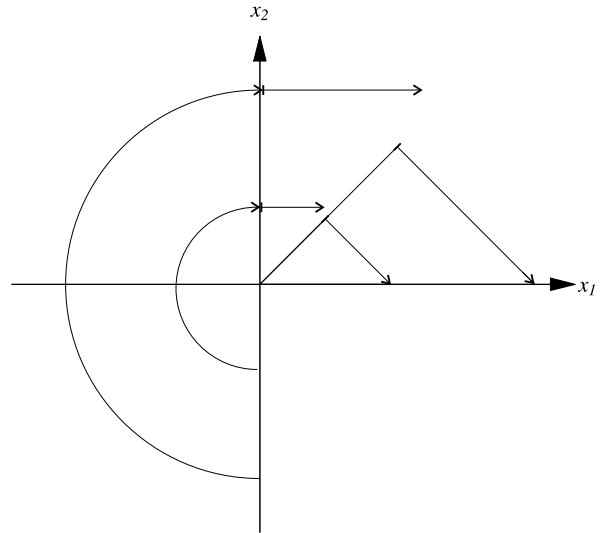


Abbildung 7.24: Das Feld aus Beispiel 7.11.3 beschreibt die Geschwindigkeit auf einem Karussell!

1. Man malt an jeden Ort bzw. an einige repräsentative Orte $\vec{x}$ jeweils den Vektor $\vec{f}(\vec{x})$. Für ein zeitabhängiges Feld fertigt man mehrere Bilder zu verschiedenen Zeitpunkten an.
2. Man zeichnet Stromlinien bzw. Feldlinien, mit Richtungspfeilen. Feldlinien verlaufen immer tangential zu den Vektoren des Feldes. Hierdurch stellt man nur die Richtung der Feldvektoren dar, ihr Betrag geht verloren. Für zeitabhängige Felder fertigt man wieder mehrere Bilder an.

Beispiel 7.11.3. Wir betrachten das Vektorfeld

$$f : \mathbb{R}^2 \rightarrow \mathbb{R}^2 \quad \vec{f}(\vec{x}) = \begin{pmatrix} x_2 \\ -x_1 \end{pmatrix}.$$

Was stellt es wohl dar? Wir wählen im Bild 7.24 rechts Methode 1. und links Methode 2. zur Veranschaulichung.

Bemerkung 7.11.4. Stromlinien sind im allgemeinen nicht einfach zu berechnen, wenn man nur $f(x)$ kennt. Im folgenden Kapitel über dynamische Systeme werden wir die Berechnung von Stromlinien – in anderem Gewand – aber noch ausführlich behandeln!

7.11.1 Vektorielles Kurvenintegral und Potential

Wir hatten in den letzten der Beispiele 7.11.2 schon einige Felder kennengelernt, die als Gradientenfelder $\nabla\phi$ einer skalaren Funktion ϕ aufgefasst werden können. Solche Felder spielen in der Physik eine herausragende Rolle.

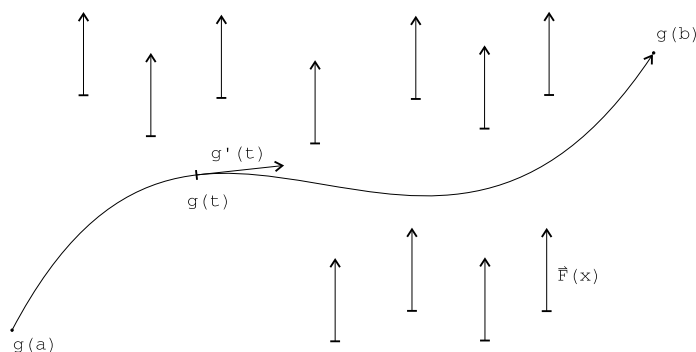


Abbildung 7.25: Ein vektoriell Kurvenintegral entlang der Kurve g im Vektorfeld $F(x)$.

Definition 7.11.5 (konservatives Vektorfeld und Potential).

Sei $f : U \rightarrow \mathbb{R}^n$, $U \subset \mathbb{R}^n$ ein Vektorfeld. Falls eine skalare Funktion $\phi : U \rightarrow \mathbb{R}$ existiert, so dass $f = \nabla\phi$, dann nennt man f ein **konservatives Vektorfeld**. Die Funktion ϕ nennt man das **Potential** von f .

Falls ϕ Potential zu einem Vektorfeld f ist, dann ist interessanterweise auch die Funktion $\tilde{\phi}(x) := \phi(x) + c$ mit beliebigem $c \in \mathbb{R}$ ein Potential zu f , denn $\nabla\tilde{\phi} = \nabla\phi = f$.

Warum man ein Feld f , für das ein Potential ϕ existiert, konservativ nennt, wird am besten klar, wenn man sich sogenannte *Kurvenintegrale* ansieht, die wir zunächst motivieren möchten.

Beispiel 7.11.6 (Energie eines Elektrons im Elektrischen Feld).

Welche Energie nimmt ein Elektron auf, das in einem elektrischen Feld bewegt wird? Im Falle einer ortsunabhängigen Kraft wissen wir, dass die Arbeit gleich Kraft mal Strecke mal dem Cosinus des Winkels zwischen Kraft und Strecke ist, und man dies mit Hilfe des Skalarproduktes • im $\mathbb{R}^3$ schreiben konnte als $W = \vec{F} \bullet \vec{s} = \|\vec{F}\| \cdot \|\vec{s}\| \cdot \cos \angle(\vec{F}, \vec{s})$, wobei $\vec{F}$ die Kraft sein soll, die auf das Elektron wirkt, und $\vec{s}$ die Strecke, die es bewegt wird.

Wir wollen nun aber den Fall zulassen, dass die Kraft $\vec{F}$ vom Ort abhängt, also ein Vektorfeld $\vec{F} : U \rightarrow \mathbb{R}^3$, $\vec{x} \mapsto \vec{F}(\vec{x})$ ist ($U \subset \mathbb{R}^3$). Ausserdem wollen wir uns nicht auf einer geraden, sondern auf einer gekrümmten Strecke bewegen, die durch eine Kurve $\vec{g} \in C^1([a, b], U)$ beschrieben

wird. Wir können bereits die Gesamtlänge $\int_a^b \|\vec{g}'(t)\| dt$ einer solchen Kurve berechnen (siehe Definition 7.1.6), aber das hilft uns beim Ermitteln der aufgenommenen Energie nicht, denn die Kraft $\vec{F}$ ist im Allgemeinen nicht an allen Punkten gleich.

Stattdessen müssen wir *infinitesimale* Wegstücke $d\vec{s}$ auf der Kurve betrachten, das sind unendlich kurze Vektoren, die kleine Stücke der Kurve g repräsentieren. Dann könnten wir das bisher nicht ganz sauber definierte Integral

$$W = \int_g \vec{F}(\vec{x}) \bullet d\vec{s}$$

zur Berechnung der Arbeit verwenden, siehe Abbildung 7.25. Durch geometrische Einsicht können wir uns davon überzeugen, dass der Ausdruck $d\vec{s} := d\vec{g} = \vec{g}'(t)dt$ gerade das gewünschte Streckenelement an einer Stelle t auf der Kurve liefert ($\vec{g}'(t)$ ist ja gerade der Tangentialvektor).

Wir berechnen W also konkret als

$$W = \int_a^b \vec{F}(\vec{g}(t)) \bullet \vec{g}'(t) dt, \quad \text{und dies kann man auch wie folgt interpretieren:}$$

$$W = \int_a^b \underbrace{\vec{F}(\vec{g}(t))}_{\text{Kraft}} \bullet \underbrace{\frac{\vec{g}'(t)}{\|\vec{g}'(t)\|}}_{\cos \angle(\vec{F}, \vec{g}')} \cdot \underbrace{\|\vec{g}'(t)\|}_{\text{Länge}} dt.$$

Definition 7.11.7 (vektorielles Kurvenintegral).

Sei $f : U \rightarrow \mathbb{R}^n$, $U \subset \mathbb{R}^n$ ein Vektorfeld und $g \in C^1([a, b], U)$ eine stetig differenzierbare Kurve. Das **vektorielle Kurvenintegral** oder **Wegintegral** ist definiert als

$$\int_a^b \langle f(g(t)), g'(t) \rangle dt.$$

Wir können uns jetzt fragen, welchen Wert das Kurvenintegral annimmt, wenn man über verschiedene Kurven $g, \tilde{g}$ vom gleichen Anfangspunkt $g(a) = \tilde{g}(\tilde{a})$ zum gleichen Endpunkt $g(b) = \tilde{g}(\tilde{b})$ geht. Es gilt folgender erstaunlicher Satz:

Satz 7.11.8 (Unabhängigkeit des Kurvenintegrals vom Weg bei konservativen Feldern).

Sei $f : U \rightarrow \mathbb{R}^n$, $U \subset \mathbb{R}^n$ ein Vektorfeld. f ist genau dann konservativ, wenn für zwei beliebige Kurven $g \in C^1([a, b], U)$, $\tilde{g} \in C^1([\tilde{a}, \tilde{b}], U)$ mit $g(a) = \tilde{g}(\tilde{a})$ und $g(b) = \tilde{g}(\tilde{b})$ gilt, dass ihre vektoriellen Kurvenintegrale gleich sind.

$$\int_a^b \langle f(g(t)), g'(t) \rangle dt = \int_{\tilde{a}}^{\tilde{b}} \langle f(\tilde{g}(t)), \tilde{g}'(t) \rangle dt.$$

Wir beweisen hier nur die erste Richtung der Äquivalenz, nämlich dass die Konservativität eines Feldes die Unabhängigkeit des Kurvenintegrals vom Weg impliziert. Dafür zeigen wir zunächst folgendes Lemma

Lemma 7.11.9 (Potential und Wegintegral).

Sei $f : U \rightarrow \mathbb{R}^n$, $U \subset \mathbb{R}^n$ ein konservatives Vektorfeld, und $\phi \in C^1(U, \mathbb{R})$ ein Potential für f . Dann gilt:

$$\int_a^b \langle f(g(t)), g'(t) \rangle dt = \phi(g(b)) - \phi(g(a)). \quad (7.12)$$

Aus dem Lemma folgt sofort die erste Richtung der Äquivalenz in Satz 7.11.8, denn

$$\phi(g(b)) - \phi(g(a)) = \phi(\tilde{g}(\tilde{b})) - \phi(\tilde{g}(\tilde{a})).$$

Um Formel (7.12) im Lemma zu beweisen, nutzen wir die Kettenregel für Jacobi-Matrizen aus Satz 7.6.6. Denn es gilt

$$\begin{aligned}
 \int_a^b \langle f(g(t)), g'(t) \rangle dt &= \int_a^b \langle \nabla \phi(g(t)), g'(t) \rangle dt \\
 &= \int_a^b \nabla \phi(g(t))^T \cdot g'(t) dt \\
 &= \int_a^b \phi'(g(t)) \cdot g'(t) dt \\
 &= \int_a^b (\phi \circ g)'(t) dt \\
 &= (\phi \circ g)(t) \Big|_a^b \\
 &= \phi(g(b)) - \phi(g(a)).
 \end{aligned}$$

□

Zusammenhang Wegunabhängigkeit und vollständiges Differential

$$\int_C \vec{F} \cdot d\vec{r} = \int_C F_x dx + F_y dy$$

$F_x dx + F_y dy$ ist vollständiges Differential $\Leftrightarrow \exists$ „Potential“ Φ mit $\nabla \Phi = \begin{pmatrix} F_x \\ F_y \end{pmatrix}$, d.h.

$$\frac{\partial \Phi}{\partial x} = F_x, \quad \frac{\partial \Phi}{\partial y} = F_y, \quad d\Phi = F_x dx + F_y dy$$

$$\Leftrightarrow \int_C \vec{F} d\vec{r} = \int_C \vec{F} \cdot r'(t) dt$$

ist wegunabhängig (konservativ). Berechne dann

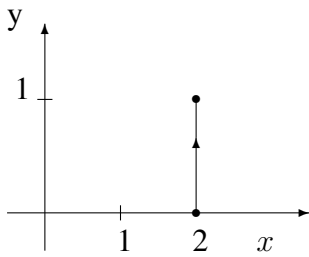
$$\int_C \vec{F} d\vec{r} = \int_C F_x dx + F_y dy = \int_C d\Phi = \int_{P_1}^{P_2} d\Phi = \Phi(P_2) - \Phi(P_1),$$

mit P_1 Anfangspunkt und P_2 Endpunkt der Kurve C .

Beispiel 7.11.10. Berechne $\int_C 2xy^3 dx + 3x^2 y^2 dy$

Kurve (Weg) $C : (0, 0) \rightarrow (2, 0) \rightarrow (2, 1)$ in der Koordinatenebene.

$$\int_C 2xy^3 dx + 3x^2 y^2 dy = \int_{(x,y)=(0,0)}^{(x,y)=(2,1)} 2xy^3 dx + 3x^2 y^2 dy$$



$$\begin{aligned}
 &= \int_{(0,0)}^{(2,0)} \underbrace{2xy^3 dx + 3x^2 y^2 dy}_{=0} + \int_{(2,0)}^{(2,1)} \underbrace{2xy^3 dx + 3x^2 y^2 dy}_{x=2=\text{const. einsetzen}} \\
 &\quad \int_{y=0}^{y=1} 3 \cdot 4y^2 dy = [4y^3]_0^1 = 4
 \end{aligned}$$

$2xy^3 dy + 3x^2 y^2 dy$ ist vollständiges Differential (s. Kap.7.5), d.h. nach 7.9.8 folgt

$$\int_C 2xy^3 dy + 3x^2 y^2 dy = \int_{P_1}^{P_2} d\Phi$$

($P_1 = (0, 0)$: Anfangspunkt von C , $P_2 = (2, 1)$: Endpunkt von C)

$$= \Phi(P_2) - \Phi(P_1) = 2^2 \cdot 1^3 - 0^2 \cdot 0^3 = 4$$

mit $\Phi = x^2 y^3$, denn

$$d\Phi = \frac{\partial \Phi}{\partial x} dx + \frac{\partial \Phi}{\partial y} dy = 2xy^3 dx + 3x^2 y^2 dy$$

Das Integral ist wegunabhängig!

Das Linienintegral über ein vollständiges Differential ist immer wegunabhängig! → wichtige Anwendungen in der Physikalischen Chemie (Thermodynamik)

7.11.2 Quellen und Senken von Vektorfeldern

Eine wichtige Frage ist: Wo entstehen, wo verschwinden Stromlinien eines Vektorfeldes? Wir motivieren dies an einem Beispiel: Auf Wetterkarten sind die Windrichtungen an einigen Orten durch Pfeile dargestellt, und dazu sind oft Isobaren-Linien angegeben, das sind Niveaumengen konstanten Drucks, siehe Abbildung 7.26. Wenn man so eine Karte genau betrachtet, sieht man, dass die Windrichtungspfeile die Isobaren fast immer schneiden, und zwar von höherem zu niedrigerem Druck. Daraus schließen wir, dass die Stromlinien des Windes in der Summe aus Hochdruckgebieten herauslaufen, und in Tiefdruckgebiete hinein.

Wie kann man nun mathematisch Orte charakterisieren, aus denen mehr Stromlinien herauslaufen als hinein, oder umgekehrt? Mathematisch wird dies durch den Begriff der *Divergenz* beschrieben.

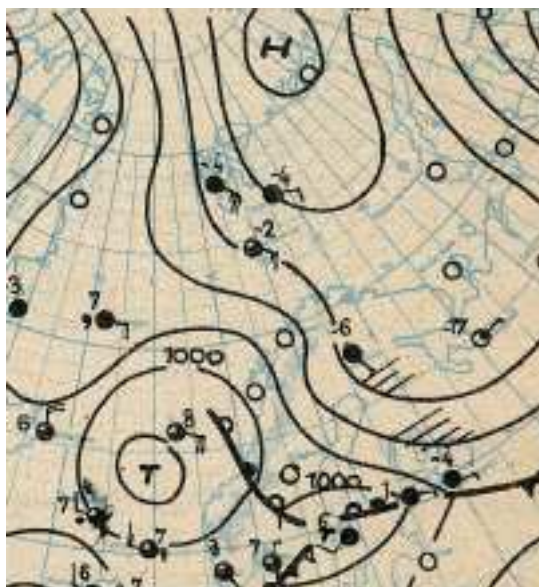


Abbildung 7.26: Wetterkarte. Der Wind weht aus dem Hoch- ins Tiefdruckgebiet.

Definition 7.11.11 (Divergenz, Quellen, Senken).

Sei $f \in C^1(U, \mathbb{R}^3)$, $U \subset \mathbb{R}^3$, dann ist die **Divergenz** von f gegeben durch

$$\operatorname{div} \vec{f}(\vec{x}) := \frac{\partial f_1}{\partial x_1}(x) + \frac{\partial f_2}{\partial x_2}(x) + \frac{\partial f_3}{\partial x_3}(x).$$

Man schreibt oft auch $\langle \nabla, f \rangle$ oder $\nabla \bullet f$ statt $\operatorname{div} f$ (wobei $\bullet$ ein Skalarprodukt symbolisieren soll). Man bezeichnet Orte, an denen die Divergenz positiv ist, als **Quellen**, und Orte mit negativer Divergenz als **Senken** von f .

Auf der Wetterkarte sind also die Hochdruckgebiete die Quellen des Windfeldes, und die Tiefdruckgebiete die Senken.

Anschauliche Motivation und Bedeutung der Divergenz am Beispiel eines Störungsfelds:

$$\begin{pmatrix} x \\ y \\ z \end{pmatrix} \mapsto \begin{pmatrix} v_x \\ v_y \\ v_z \end{pmatrix} =: v \text{ (Geschwindigkeit der Strömung an der Stelle } (x, y, z))$$

$$\begin{array}{c} \text{Einstrom} \\ \xrightarrow{v_y(x, y, z)} \end{array} \Delta z \begin{array}{c} \Delta x \\ \Delta y \end{array} \xrightarrow{\text{Ausstrom}} \begin{array}{c} v_y(x, y + \Delta y, z) \end{array}$$

Flüssigkeitsmenge pro Zeitintervall Δt :

$$\text{Einstrom } v_y(x, y, z) \cdot \Delta x \Delta z \Delta t$$

Ausstrom $v_y(x, y + \Delta y, z) \cdot \Delta x \Delta z \Delta t$

Differenz pro Zeiteinheit:

$$[v_y(x, y + \Delta y, z) - v_y(x, y, z)] \cdot \Delta x \Delta z = \frac{[v_y(x, y + \Delta y, z) - v_y(x, y, z)]}{\Delta y} \cdot \underbrace{\Delta x \Delta y \Delta z}_{=\Delta v}$$

d.h. pro Volumeneinheit: Differenz Ausstrom-Einstrom pro Zeit- und Volumeneinheit =

$$\frac{v_y(x, y + \Delta y, z) - v_y(x, y, z)}{\Delta y}, \text{ im Grenzwert } \Delta y \rightarrow 0 \text{ also } \rightarrow \frac{\partial v_y}{\partial y}$$

analog für x - und z -Richtung: $\frac{\partial v_x}{\partial x}, \frac{\partial v_z}{\partial z}$ und insgesamt (Summe):

$$\frac{\partial v_x}{\partial x} + \frac{\partial v_y}{\partial y} + \frac{\partial v_z}{\partial z} =: \operatorname{div} v$$

„Volumengewinn“ (Differenz Ausstrom - Einstrom) pro Zeit- und Volumeneinheit

Beispiel 7.11.12. Wir geben einige Beispiele zur Illustration von Quellen und Senken.

$$1. \vec{f}(\vec{x}) = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix}, \quad \operatorname{div} \vec{f} = 1 + 1 + 1 = 3.$$

Bei diesem Feld sind alle Orte Quellen.

$$2. \vec{f}(\vec{x}) = \begin{pmatrix} x_2 \\ -x_1 \\ 0 \end{pmatrix}, \quad \operatorname{div} \vec{f} = 0 + 0 + 0 = 0.$$

Dieses Feld beschreibt eine einfache Drehung, und nirgendwo sind Quellen oder Senken.

$$3. \vec{f}(\vec{x}) = \begin{pmatrix} \sin x_1 \\ 0 \\ 0 \end{pmatrix}, \quad \operatorname{div} \vec{f} = \cos x_1 + 0 + 0 = \cos x_1.$$

Dieses Feld beschreibt z.B. das Luftgeschwindigkeitsfeld bei einer (longitudinalen) Schallwelle in x_1 -Richtung (zu einem festen Zeitpunkt). Quellen und Senken wechseln sich in x_1 -Richtung ab.

$$4. \vec{f}(\vec{x}) = \begin{pmatrix} 0 \\ \sin x_1 \\ 0 \end{pmatrix}, \quad \operatorname{div} \vec{f} = 0 + 0 + 0 = 0.$$

Dieses Feld beschreibt z.B. das Geschwindigkeitsfeld einer (transversalen) Scherwelle in x_1 -Richtung, die in x_2 -Richtung schwingt. Es gibt keine Quellen und Senken.

Man charakterisiert Quellen und Senken eines Vektorfeldes über die Divergenz:

$\operatorname{div} v > 0$: Volumenelement ΔV enthält eine (sog.) „Quelle“, d.h. es strömt mehr Flüssigkeit aus als ein

$\operatorname{div} v < 0$: ΔV enthält eine „Senke“

$\operatorname{div} v = 0$: Es gilt ein Erhaltungsgesetz in ΔV (Ausstrom = Einstrom)

Verallgemeinerung auf beliebiges Vektorfeld möglich: ersetze v durch andere vektorielle physikalische Größen (z.B. elektr./magnet. Feld)

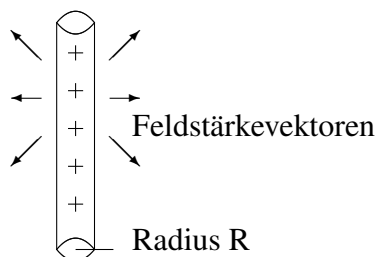
Bemerkung 7.11.13. Die Divergenz eines Vektorfeldes ist selbst ein Skalarfeld !

Beispiel 7.11.14. Elektrisches Feld eines geladenen langen Zylinders

$$\vec{E} = \frac{\varphi R^2}{2\epsilon_0(x^2 + y^2)} \cdot \begin{pmatrix} x \\ y \\ 0 \end{pmatrix}$$

(elektr. Feld außerhalb des Zylinders)

Berechne die Divergenz im Punkt (x, y) :



$$\frac{\partial E_x}{\partial x} = \frac{\partial R^2}{2\epsilon_0} \frac{\partial}{\partial x} \left(\frac{x}{x^2 + y^2} \right)$$

Quotientenregel

$$\begin{aligned} &= \frac{\varphi R^2}{2\epsilon_0} \cdot \frac{1 \cdot (x^2 + y^2) - x \cdot 2x}{(x^2 + y^2)^2} = \frac{\varphi R^2}{2\epsilon_0} \cdot \frac{y^2 - x^2}{(x^2 + y^2)^2} \\ \frac{\partial E_y}{\partial y} &= \frac{\varphi R^2}{2\epsilon_0} \cdot \frac{\partial}{\partial y} \left(\frac{y}{x^2 + y^2} \right) = \frac{\varphi R^2}{2\epsilon_0} \cdot \frac{x^2 - y^2}{(x^2 + y^2)^2} \\ \frac{\partial E_z}{\partial z} &= 0, \operatorname{div} \vec{E} = \frac{\partial E_x}{\partial x} + \frac{\partial E_y}{\partial y} + \frac{\partial E_z}{\partial z} = 0 \end{aligned}$$

→ elektr. Feld außerhalb des Zylinders ist quellenfrei! Elektrisches Feld innerhalb des Zylinders:

$$\vec{E} = \frac{\varphi}{2\epsilon_0} \begin{pmatrix} x \\ y \\ 0 \end{pmatrix}$$

$$\operatorname{div} \vec{E} = \frac{\partial E_x}{\partial x} + \frac{\partial E_y}{\partial y} + \frac{\partial E_z}{\partial z} = \frac{\varphi}{2\epsilon_0} + \frac{\varphi}{2\epsilon_0} + 0 = \frac{\varphi}{\epsilon_0}$$

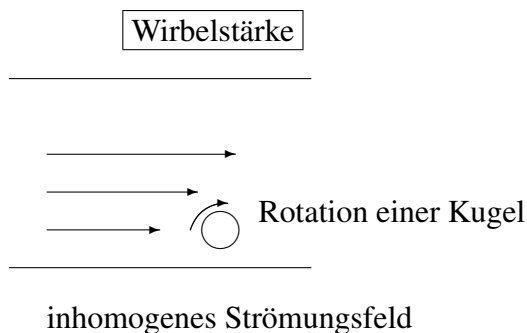
(φ ist die Ladungsdichte und $\operatorname{div} \vec{E} = \frac{\varphi}{\epsilon_0}$ ist eine der Maxwell'schen Grundgleichungen der Elektrodynamik in der Physik.)

Definition 7.11.15 (Rotation eines Vektorfelds). Sei $v : U \rightarrow \mathbb{R}^3$ ein Vektorfeld. Definiere den Differentialoperator, $U \subset \mathbb{R}^3$

$$\operatorname{rot} v := \begin{pmatrix} \frac{\partial v_z}{\partial y} - \frac{\partial v_y}{\partial z} \\ \frac{\partial v_x}{\partial z} - \frac{\partial v_z}{\partial x} \\ \frac{\partial v_y}{\partial x} - \frac{\partial v_x}{\partial y} \end{pmatrix} = \nabla \times v \text{ (Vektorprodukt)}$$

$$\begin{pmatrix} \frac{\partial}{\partial x} \\ \frac{\partial}{\partial y} \\ \frac{\partial}{\partial z} \end{pmatrix} \times \begin{pmatrix} v_x \\ v_y \\ v_z \end{pmatrix}$$

Bemerkung 7.11.16. rot ist nur für Vektorfelder im $\mathbb{R}^3$ definiert. $\operatorname{rot} v$ ist selbst wieder ein Vektorfeld (auch Wirbelfeld genannt).



Anschauliche Interpretation der Rotation: Richtung von $\operatorname{rot} v$ entspricht der Drehachse einer hypothetischen Kugel, der Betrag (Norm) von $\operatorname{rot} v$ ist Maß für die sog. „Wirbelstärke“ („Geschwindigkeit der Rotation“)

Es gelten folgende Relationen:

1. $f = \operatorname{rot} E$ (E Vektorfeld) $\Rightarrow \operatorname{div} f = 0$ (Wirbelfelder sind quellenfrei)
2. $\operatorname{div} f = 0 \Rightarrow$ Es gibt ein Vektorfeld E mit $f = \operatorname{rot} E$ (Ein quellenfreies Vektorfeld lässt sich als Wirbelfeld darstellen)
3. $\operatorname{rot} (\operatorname{grad} \Phi) = \operatorname{rot} (\nabla \Phi) = 0$ (Φ Skalarfeld) (Ein Gradientenfeld ist wirbelfrei)
4. $\operatorname{rot} f = 0 \Rightarrow \exists$ Skalarfeld Φ mit $f = \operatorname{grad} \Phi = \nabla \Phi$

Notation mit Hilfe des Nabla-Operators im $\mathbb{R}^3$

$$\nabla := \begin{pmatrix} \frac{\partial}{\partial x} \\ \frac{\partial}{\partial y} \\ \frac{\partial}{\partial z} \end{pmatrix}$$

$$\operatorname{div} f = \langle \nabla, f \rangle$$

$$\operatorname{rot} f = \nabla \times f \text{ (Vektorprodukt)}$$

$$\operatorname{div}(\operatorname{grad} f) = \langle \nabla, \nabla \rangle = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2} =: \Delta \text{ (Laplace-Operator)}$$

$$\Delta f = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} + \frac{\partial^2 f}{\partial z^2}$$

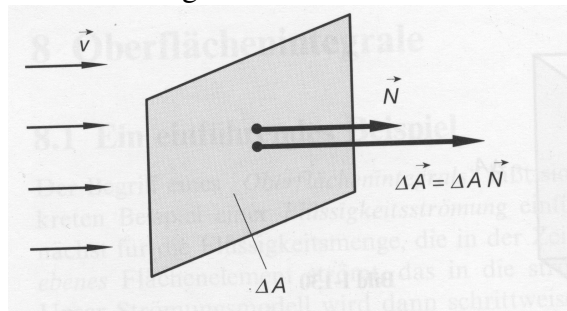
Alle diese Vektoroperatoren sind Differentialoperatoren (Abbildungen, die Funktionen spezielle Formen ihrer Ableitungen zuordnen)

Relation 4. ist von großer Bedeutung und stellt eine hinreichende Bedingung für die Existenz einer „Stammfunktion“ eines Vektorfeldes dar, in dem Sinne, dass das Vektorfeld das Gradientenfeld eines geeigneten Skalarfeldes ist. Das motiviert die Definitionen **konservatives Vektorfeld und Potential**.

Falls Φ Potential zu einem Vektorfeld f ist, dann ist auch die Funktion $\Phi(x) := \Phi(x) + c$ mit beliebigem $c \in \mathbb{R}$ ein Potential zu f , denn $\nabla \Phi = \nabla \Phi = f$.

7.11.3 Flächenelemente von Oberflächen

Die Abbildungen in diesem Abschnitt stammen aus [Pap].



$\vec{N}$: normierter Normalenvektor ($\vec{N}$ orthogonal zu A):

$$||\vec{N}|| = 1$$

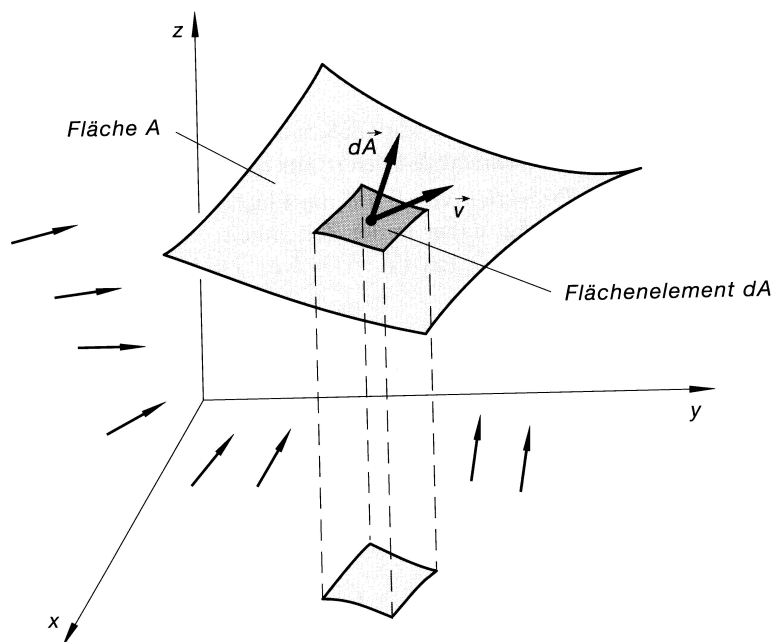
$$||\Delta \vec{A}|| = \Delta A \underbrace{||\vec{N}||}_1 = \Delta A$$

Definiere einen „Flächenelementvektor“ $\Delta \vec{A}$ als Normalenvektor mit der Länge (der Norm) ΔA (Flächeninhalt).

Bei gekrümmten Oberflächen:

$$\lim_{\Delta \rightarrow 0} \Delta A = dA$$

$$d\vec{A} = \vec{N} dA$$

Definition Oberflächenintegral

Oberflächenintegral:

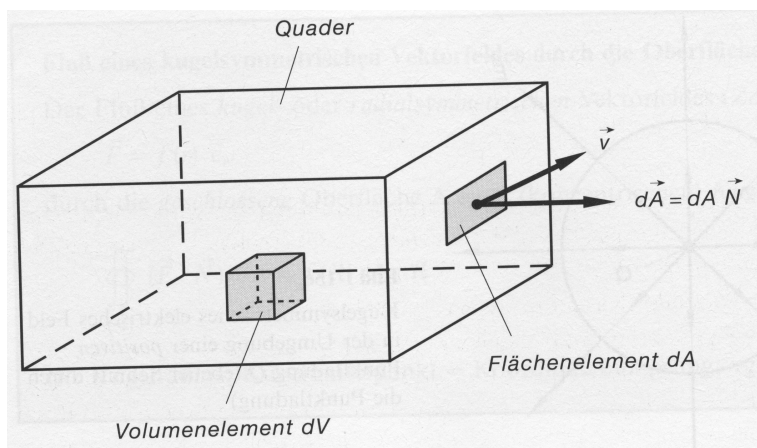
$$\iint_A \vec{v} d\vec{A} := \iint_A \langle \vec{v}, \vec{N} \rangle dA$$

7.11.4 Der Gauß'sche Integralsatz

Das Oberflächenintegral eines Vektorfeldes $\vec{f}$ über eine **geschlossene** Fläche A ist gleich dem Volumenintegral der Divergenz $\text{div } \vec{f}$ von $\vec{f}$ über das von A eingeschlossene Volumen.

$$\boxed{\oiint_A \vec{f} \cdot d\vec{A} = \oiint_A \langle \vec{f}, \vec{N} \rangle dA = \iiint_V \text{div } \vec{f} dV}$$

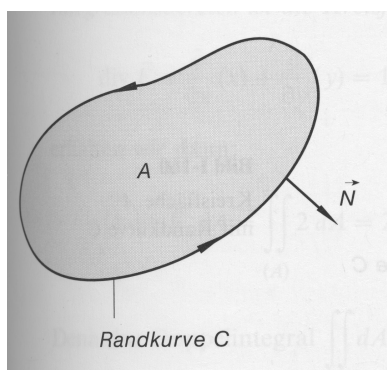
Anschauliche Interpretation: Oberflächenintegral beschreibt die “Fluss”-Bilanz (Differenz Ausstrom - Einstrom) einer physikalischen Größe durch die Oberfläche A .



Das Volumenintegral $\text{ber } \text{div } \vec{f}$ beschreibt den “Gewinn” der physikalischen Größe im eingeschlossenen Volumen V

Satz von Gauss in der Ebene:

$$\oint_C \vec{f} d\vec{s} = \oint \langle \vec{f}, \vec{N} \rangle ds = \iint_A \text{div } \vec{f} dA$$



7.11.5 Satz von Stokes

Satz 7.11.17 (Satz von Stokes). Das Kurven(linien-)Integral eines Vektorfeldes $\vec{f}$ entlang einer geschlossenen Kurve C ist gleich dem Oberflächenintegral der Rotation $\text{rot } \vec{f}$ von $\vec{f}$ ber die von C berandete Fläche A .

$$\oint_C \vec{f} \cdot d\vec{r} = \iint \text{rot } \vec{f} d\vec{A} = \iint \langle \text{rot } \vec{f}, \vec{N} \rangle dA$$

Die Sätze von Gauss und Stokes gehören zu den wichtigsten Sätzen der gesamten Analysis und stellen Verallgemeinerungen des Hauptsatzes der Differential-und Integralrechnung dar.

Kapitel 8

Gewöhnliche Differentialgleichungen und Dynamische Systeme

Spezialliteratur (Dynamik) mit Fokus auf biologischen Anwendungen

S. H. Strogatz	Nonlinear dynamics and chaos, with applications to physics, biology, chemistry and engineering
A. Goldbeter	Biochemical oscillations and cellular rhythms
C. Fall et al.	Computational Cell Biology
J. Keener, J. Sneyd	Mathematical Physiology
J. D. Murray	Mathematical Biology
L. Segel	Modeling Dynamic Phenomena in Molecular and Cellular Biology
F. W. Schneider	Nichtlineare Dynamik in der Chemie
F. Verhulst	Nonlinear differential equations and dynamical systems, (sehr mathematische Darstellung)
V. I. Arnold	Gewöhnliche Differentialgleichungen (mathematische Darstellung)
E. Klipp et al.	Systems Biology in Practice

Kontinuierliche dynamische Systeme

Definition 8.0.18 (Dynamisches System). Ein metrischer Raum X (= „Menge X mit einer Norm „=“ Maß für Abstände“) und eine stetige Abbildung $\times [0, \infty) \rightarrow X$ mit

$$F(x, 0) = x, \quad F(F(x, s), t) = F(x, s + t)$$

Wir betrachten konkret zeitabhängige Differentialgleichungen. X ist in der Regel ein Vektorraum, meistens der $\mathbb{R}^n$.

$\dot{x} = \frac{dx}{dt} = f(x, t)$ gewöhnliche Differentialgleichung (ODE): f =Funktion; partielle Differentialgleichung (PDE): f =Operator. Man nennt die Gleichung auch Evolutionsgleichung.

Definition 8.0.19. Eine gewöhnliche Differentialgleichung n -ter Ordnung ist eine Gleichung einer Funktion $x(t)$, in der Ableitungen von x nach t bis zur Ordnung n vorkommen.

Bemerkung 8.0.20. Wir werden uns i.W. auf Funktion $x(t)$ beschränken, in denen die unabhängige Variable die Zeit (t) darstellt. Die Definition gilt aber für beliebige Funktionen.

Beispiel 8.0.21.

$$x + y(x) \cdot \underbrace{y'(x)}_{=\frac{dy}{dx}} = 0$$

Die allgemeine Lösung einer Differentialgleichung n -ter Ordnung enthält n freie Parameter. Man erhält die Lösung durch Integration der Differentialgleichung und die freien Parameter entsprechen den Integrationskonstanten.

Anwendungen:

Man ist oft an speziellen Lösungen interessiert. Diese werden durch zusätzliche Anfangs- oder Randbedingungen festgelegt.

Beispiel 8.0.22. $y'(x) = 2x$, $y(0) = 1$ (Anfangswertproblem)

$$\int y'(x) dx = \int 2x dx \Leftrightarrow y(x) = x^2 + C$$

$$y(0) = 0^2 + C = 1 \Leftrightarrow C = 1$$

→ spezielle Lösung $y(x) = x^2 + 1$

Beispiel 8.0.23.

$$y''(x) = -\frac{1}{2}(3x - x^2), \quad y(0) = y(3) = 0, \quad x \in [0, 3]$$

$$y'(x) = \int y''(x) dx = -\frac{1}{2} \int (3x - x^2) dx = -\frac{1}{2} \cdot \left(\frac{3}{2}x^2 - \frac{1}{3}x^3 + C_1 \right)$$

$$y(x) = \int y'(x) dx = -\frac{1}{2} \int \left(\frac{3}{2}x^2 - \frac{1}{3}x^3 + C_1 \right) dx = -\frac{1}{2} \cdot \left(\frac{1}{2}x^3 - \frac{1}{12}x^4 + C_1x + C_2 \right)$$

$$y(0) = 0 \Leftrightarrow -\frac{1}{2} \left(\frac{1}{2}x^3 - \frac{1}{12}x^4 + C_1x + C_2 \right) = 0 \Leftrightarrow C_2 = 0$$

$$y(3) = 0 \Leftrightarrow -\frac{1}{2} \left(\frac{1}{2} \cdot 3^3 - \frac{1}{12}3^4 + C_1 \cdot 3 \right) = -\frac{1}{2} \left(\frac{27}{2} - \frac{81}{12} + C_1 \cdot 3 \right) = 0 \Leftrightarrow C_1 = \frac{9}{4}$$

Spezielle Lösung: $y(x) = -\frac{1}{2} \left(\frac{1}{2}x^3 - \frac{1}{12}x^4 + \frac{9}{4}x \right)$

Merke: Zur Bestimmung von n Integrationskonstanten sind n Bedingungen in Form von Anfangs- oder Randbedingungen nötig! → führt i.a. auf ein n -dim. lineares Gleichungssystem.

Definition 8.0.24. Eine Differentialgleichung 1.Ordnung heißt linear, wenn man sie in der Form $y'(x) + f(x) \cdot y(x) = g(x)$ schreiben kann. Man nennt sie homogen, wenn $g(x) = 0$, anderenfalls inhomogen.

Lösungsverfahren:

- homogene, lineare Differentialgleichung 1. Ordnung
→ Separation der Variablen
- inhomogen
→ Lösung des homogenen Falls mit $g(x) = 0$, dann “Variation der Konstanten”

in allgemeiner Form:

$$y' + f(x)y = g(x), \text{ Setze zunchst } g(x) = 0$$

$$\longrightarrow \frac{dy}{dx} + f(x)y = 0$$

Trennung der Variablen $\longrightarrow \frac{dy}{y} = -f(x)dx$ Integration:

$$\int \frac{dy}{y} = \int -f(x)dx$$

$$\ln y = - \int f(x)dx + C$$

$$y = K \cdot e^{-\int f(x)dx}, K = e^C$$

(allgemeine Lösung der homogenen Gleichung)

Ersetze Integrationskonstante K durch eine zunchst unbekannte Funktion K(x) und bestimme K(x) durch **Variation der Konstanten**: $y = K(x)e^{-\int f(x)dx}$

Einsetzen in Differentialgleichung $y' + f(x)y = g(x)$:

$$y' = K'(x)e^{-\int f(x)dx} - K(x)f(x)e^{-\int f(x)dx}$$

(Produktregel)

$$\overbrace{K'(x)e^{-\int f(x)dx} - K(x)f(x)e^{-\int f(x)dx}}^{y'} + \underbrace{f(x)K(x)e^{-\int f(x)dx}}_{f(x)y} = g(x)$$

$$\Leftrightarrow K'(x)e^{-\int f(x)dx} = g(x)$$

$$\Leftrightarrow K'(x)g e^{\int f(x)dx}$$

$$\rightarrow K() = \int g(x)e^{\int f(x)dx} + C$$

→ allgemeine Lösung inhomogener Fall:

$$y(x) = K(x)e^{-\int f(x)dx} = \left(\int g(x)e^{\int f(x)dx} + C \right) e^{-\int f(x)dx}$$

Lineare Systeme gewöhnlicher Differentialgleichungen

$$\dot{x}(t) = \frac{dx}{dt} = f(t, x), x(t) \in \mathbb{R}^n$$

↓

lineare Funktion $f : \mathbb{R} \times \mathbb{R}^n \rightarrow \mathbb{R}^n, (t, x) \mapsto f(t, x)$

$$\begin{pmatrix} \dot{x}_1(t) \\ \dot{x}_2(t) \\ \vdots \\ \dot{x}_n(t) \end{pmatrix} = \begin{pmatrix} f_1(t, x) \\ f_2(t, x) \\ \vdots \\ f_n(t, x) \end{pmatrix}$$

System gewöhnlicher Differentialgleichungen

Allgemeines Lösungsverfahren für lineare Systeme gewöhnlicher Differentialgleichungen (Entkopplung). am Beispiel

$$\dot{x}_1(t) = 6x_1(t) + 2x_2(t)$$

$$\dot{x}_2(t) = 2x_1(t) + 3x_2(t)$$

in Matrixschreibweise:

$$\dot{x} = \begin{pmatrix} \dot{x}_1 \\ \dot{x}_2 \end{pmatrix} = Ax = \begin{pmatrix} 6 & 2 \\ 2 & 3 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$$

Idee der Entkopplung:

Diagonalisierung von A !

Eigenwerte:

$$|A - \lambda I| = 0 = \begin{vmatrix} 6 - \lambda & 2 \\ 2 & 3 - \lambda \end{vmatrix} = 0$$

$$\Leftrightarrow \lambda^2 - 9\lambda + 14 = 0$$

$$\lambda_{1,2} = \frac{9}{2} \pm \sqrt{\frac{81}{4} - \frac{56}{4}}$$

$$= \frac{9}{2} \pm \sqrt{\frac{25}{4}}$$

$$= \frac{9}{2} \pm \frac{5}{2} \longrightarrow \lambda_1 = 7, \lambda_2 = 2$$

Eigenvektorberechnung:

$$\begin{aligned}
 (A - \lambda I)v &= 0 \\
 1. \lambda_1 = 7 \begin{pmatrix} 6-7 & 2 \\ 2 & 3-7 \end{pmatrix} \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} &= 0 \\
 \Rightarrow -v_1 + 2v_2 &= 0 \\
 \text{wähle } v_1 = 2 &\Rightarrow v_2 = 1
 \end{aligned}$$

1. Eigenvektor (zu λ_1): $v = \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$ (ein spezieller Eigenvektor zum Eigenwert $\lambda_1 = 7$)

2. $\lambda_2 = 2$

$$\begin{aligned}
 \begin{pmatrix} 6-2 & 2 \\ 23-2 \end{pmatrix} \begin{pmatrix} w_1 \\ w_2 \end{pmatrix} &= 0 \\
 4w_1 + 2w_2 &= 0 \\
 \text{wähle } w_1 = 1 &\Rightarrow w_2 = -2 \\
 \longrightarrow \text{2. Eigenvektor: } w &= \begin{pmatrix} w_1 \\ w_2 \end{pmatrix} = \begin{pmatrix} 1 \\ -2 \end{pmatrix} \\
 B := \begin{pmatrix} v & w_1 \\ v_2 & w_2 \end{pmatrix} = \begin{pmatrix} 2 & 1 \\ 1 & -2 \end{pmatrix} &\quad \begin{array}{l} \text{Eigenvektoren von } A \text{ als} \\ \text{Spalten der Matrix } B \end{array} \\
 \begin{pmatrix} 2 & 1 \\ 1 & -2 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} & \\
 \downarrow & \\
 \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} \frac{2}{5} & \frac{1}{5} \\ \frac{1}{5} & -\frac{2}{5} \end{pmatrix} &= B^{-1}
 \end{aligned}$$

Probe: $BB^{-1} = I$

$$\begin{pmatrix} 2 & 1 \\ 1 & -2 \end{pmatrix} \begin{pmatrix} \frac{2}{5} & \frac{1}{5} \\ \frac{1}{5} & -\frac{2}{5} \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$$

Probe: Diagonalisierung $B^{-1}AB$

$$\begin{pmatrix} \frac{2}{5} & \frac{1}{5} \\ \frac{1}{5} & -\frac{2}{5} \end{pmatrix} \begin{pmatrix} 6 & 2 \\ 2 & 3 \end{pmatrix} \begin{pmatrix} 2 & 11 & -2 \end{pmatrix} = \begin{pmatrix} 7 & 0 \\ 0 & 2 \end{pmatrix}$$

Lösung von $\dot{x} = Ax$?

Transformation: $y := B^{-1}x$, $x = By$

$$\dot{y} = B^{-1}\dot{x} = B^{-1}Ax = \underbrace{B^{-1}AB}_{\Lambda \text{ Diagonalmatrix}} y$$

Konkreter Fall:

$$\dot{y} = \Lambda y = \begin{pmatrix} 7 & 0 \\ 0 & 2 \end{pmatrix} y$$

$$\frac{dy_1}{dt} = 7y_1 \quad \frac{dy_2}{dt} = 2y_2 \rightarrow \begin{aligned} y_1(t) &= Ce^{7t} \\ y_2(t) &= De^{2t} \end{aligned}$$

$$\frac{y_1}{y_1} = 7dt \quad (\text{Separation der Variablen})$$

$$x(t) = \begin{pmatrix} x_1(t) \\ x_2(t) \end{pmatrix} = B \begin{pmatrix} y_1(t) \\ y_2(t) \end{pmatrix} = \begin{pmatrix} 2Ce^{7t} + 1De^{2t} \\ 1Ce^{7t} - 2De^{2t} \end{pmatrix}$$

Falls Matrix A nicht diagonalisierbar, gibt es zumindest immer eine Transformation auf sog. Jordan-Normalform:

$$B^{-1}AB = \begin{pmatrix} \lambda_1 & 1 & 0 & \dots & 0 \\ 0 & \ddots & \ddots & & \vdots \\ \vdots & & \ddots & \ddots & 0 \\ \vdots & & & \ddots & 1 \\ 0 & \dots & \dots & 0 & \lambda_n \end{pmatrix}$$

Beispiel dynamisches System aus der Physiologie

Regulation des Zuckerstoffwechsels durch Insulin:

$I(t)$: Insulinkonzentration im Blut

$Z(t)$: Zuckerkonzentration im Blut

jeweils zur Zeit t

Mathematisches Modell (Differentialgleichung)

$$\frac{dI}{dt} = \underbrace{-\alpha \cdot I \cdot Z}_{\text{Abbau}} + \underbrace{\beta \cdot Z \cdot Z}_{\text{Stimulierte Ausschüttung}}$$

$$\frac{dZ}{dt} = \underbrace{-\gamma \cdot Z \cdot I}_{\text{Abbau}} + \underbrace{\delta}_{\text{Zufuhr aus Leber etc.}}$$

Darstellung der Lösungen:

a) „Zeitreihen“ $I(t), Z(t)$ b) Phasenraumdarstellung

Bsp. $\left. \begin{aligned} \dot{x}_1 &= x_2 \\ \dot{x}_2 &= -x_1 \end{aligned} \right\} \dot{x} = \begin{pmatrix} \dot{x}_1 \\ \dot{x}_2 \end{pmatrix} = f(x) = \begin{pmatrix} x_2 \\ -x_1 \end{pmatrix}$

$$\frac{dx_1(t)}{dt} = x_2 \quad \frac{dx_2(t)}{dt} = -x_1 \rightarrow \frac{dt}{dx_2} = -\frac{1}{x_1}$$

eliminiere t :

$$\frac{dx_1}{dx_2} = \frac{dx_1}{dt} \cdot \frac{dt}{dx_2} = -\frac{x_2}{x_1} \Leftrightarrow \underbrace{x_1 dx_1 + x_2 dx_2}_{\text{Differential ist vollständig}} = 0$$

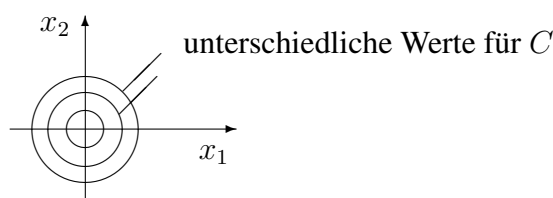
→ ∃ Potential Φ mit

$$d\Phi = x_1 dx_1 + x_2 dx_2 = 0,$$

$$\Phi = \frac{1}{2}x_1^2 + \frac{1}{2}x_2^2$$

$$d\Phi = 0 \Rightarrow \Phi = \text{const}, \text{ also } \Phi = \frac{1}{2}x_1^2 + \frac{1}{2}x_2^2 = C$$

→ die Lösungen der Differentialgleichung $\dot{x}_1 = x_2, \dot{x}_2 = -x_1$ entsprechen Kurven in der x_1, x_2 -Ebene (in diesem Fall Kreise vom Radius $\sqrt{2C}$)



Diese Abbildung nennt man „Phasenportrait“. Die Lösungskurven heißen „Trajektorien“ oder „Orbits“.

Phasenfluss

F bildet einen Anfangszustand entlang einer „Trajektorie“ auf einen Endzustand ab. Der Phasenfluss beschreibt eine Strömung im übertragenen Sinne, die durch das Richtungsfeld (das ist ein Vektorfeld!) der Differentialgleichung bestimmt ist.

In diesem Kapitel wollen wir uns der Frage widmen, wie wir die Zukunft vorhersagen können. Wir wollen dies mit Hilfe mathematischer Modelle tun, und damit die Methoden kennenlernen, die letztendlich auch bei Weltklimaprognosen, bei der Berechnung von Planetenbahnen, bei Modellen der Weltbevölkerung oder Aids-Ausbreitung, bei Wettervorhersagen usw. verwendet werden. Wir betrachten zur Motivation ein kleines Modell aus der Physiologie.

Beispiel 8.0.25 (Insulin und Blutzucker).

Wir wollen die Frage untersuchen, warum viele Menschen einige Zeit nach Verzehr eines zuckerhaltigen Müsliriegels müder werden als vor dem Verzehr. Die Mediziner erzählen uns, dass

dies daran liegt, dass durch die plötzliche Zuckerzufuhr zunächst die körpereigene Insulinausschüttung angeregt wird, und das Insulin dann zuviel Zucker abbaut, so dass am Ende für einige Zeit sogar weniger Zucker im Blut ist als zuvor: dies merkt man dann als Müdigkeit. Wir betrachten ein ganz einfaches Modell, in dem $I(t)$ die Insulinkonzentration im Blut und $Z(t)$ die von Zucker zu einem Zeitpunkt t darstellen. Wir betrachten die zeitlichen Änderungsraten $I'(t) = \frac{dI}{dt}(t)$ und $Z'(t) = \frac{dZ}{dt}(t)$, die wir durch das folgende System von *gewöhnlichen Differentialgleichungen* modellieren können:

$$\begin{aligned}\frac{dI}{dt} &= \underbrace{-\alpha \cdot I \cdot Z}_{\text{Abbau}} + \underbrace{\beta \cdot Z \cdot I}_{\text{Stimulierte Ausschüttung}} \\ \frac{dZ}{dt} &= \underbrace{-\gamma \cdot Z \cdot I}_{\text{Abbau}} + \underbrace{\delta}_{\text{Zufuhr aus Leber etc.}}\end{aligned}$$

Was passiert, wenn eine Person mit viel Zucker im Blut und wenig Insulin startet – z.B. nach Verzehr eines Müsliriegels? In Abbildung 8.1 sehen wir das Ergebnis einer Computersimulation für den Anfangswert $I(0) = 1$ und $Z(0) = 10$ (mit $\alpha = \beta = \gamma = \delta = 1$)

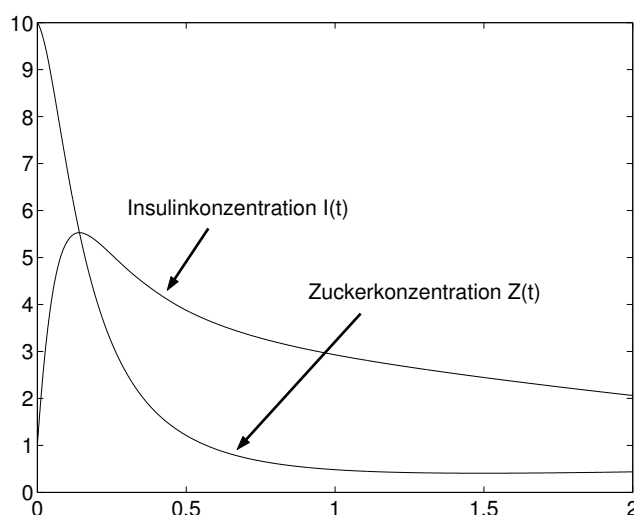


Abbildung 8.1: Insulin- und Zuckerkonzentration nach Erhöhung des Blutzuckerspiegels.

Definition 8.0.26 (gewöhnliche Differentialgleichung, Zustandsvektor).

Sei $U \subset \mathbb{R}^n$ offen, $f : [t_0, t_f] \times U \rightarrow \mathbb{R}^n$ ein zeitabhängiges Vektorfeld. Das System von n Differentialgleichungen

$$x'(t) = f(t, x(t)) \quad (8.1)$$

heißt **gewöhnliche Differentialgleichung** oder **dynamisches System**. Die Menge U heißt **Zustandsraum**, und der Vektor $x \in U$ heißt **Zustandsvektor** des Systems. Statt $x'(t)$ schreibt man oft auch $\frac{dx}{dt}(t)$ und sehr häufig auch $\dot{x}(t)$.

Definition 8.0.27 (Anfangswertproblem).

Sei $U \subset \mathbb{R}^n$ offen, $f : [t_0, t_f] \times U \rightarrow \mathbb{R}^n$ und sei $x_0 \in U$ der **Anfangswert** oder **Anfangszustand**.

Die Aufgabe:

Finde eine Kurve $x \in C^1([t_0, t_f], U)$ so dass

$$\dot{x}(t) = f(t, x(t)), \quad t \in [t_0, t_f], \quad (8.2)$$

$$\text{und} \quad x(t_0) = x_0 \quad (8.3)$$

heisst **Anfangswertproblem** (AWP). Eine Lösungskurve $x \in C^1([t_0, t_f], U)$, die (8.2) und (8.3) erfüllt, heisst **Lösung des Anfangswertproblems** oder auch **Trajektorie** zum Anfangswert x_0 .

Beispiel 8.0.28 (Insulin-Zucker-Modell).

$$x(t) = \begin{pmatrix} I(t) \\ Z(t) \end{pmatrix} \in \mathbb{R}^2$$

$$f(t, x) = \begin{pmatrix} -\alpha \cdot I \cdot Z + \beta Z^2 \\ -\gamma Z \cdot I + \delta \end{pmatrix} = \begin{pmatrix} -\alpha x_1 \cdot x_2 + \beta x_2^2 \\ -\gamma x_2 \cdot x_1 + \delta \end{pmatrix}.$$

$$\text{Der Anfangswert ist } x_0 = \begin{pmatrix} I_0 \\ Z_0 \end{pmatrix}.$$

Beim Insulin-Zucker-Modell hängt die Funktion $f(t, x)$ gar nicht direkt von der Zeit ab, sondern nur vom Zustand x . Diese Eigenschaft hat einen eigenen Namen.

Definition 8.0.29 (autonomes System).

Wenn f *nicht* von der Zeit t abhängt, man also $f(x)$ schreiben kann, sagt man, das System $\dot{x}(t) = f(x(t))$ sei **autonom**.

Im Allgemeinen ist die Frage nicht leicht zu beantworten, ob ein Anfangswertproblem überhaupt eine Lösung hat, und ob diese Lösung eindeutig ist. Glücklicherweise gibt es einen grundlegenden Satz, der für fast alle praktisch vorkommenden dynamischen Systeme die Existenz und Eindeutigkeit von Lösungen eines Anfangswertproblems garantiert.

Satz 8.0.30 (Existenz und Eindeutigkeit der Lösung des AWP).

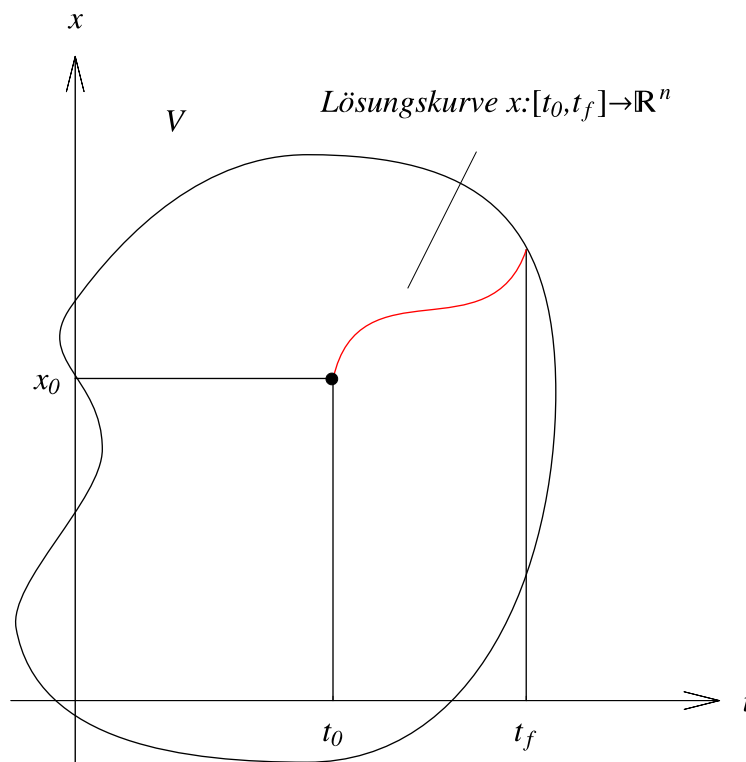
Sei $V \subset \mathbb{R} \times \mathbb{R}^n$ offen, $(t_0, x_0) \in V$, und $f \in C^1(V, \mathbb{R}^n)$ eine stetig differenzierbare Funktion. Dann gibt es ein $t_f > t_0$ so dass das AWP

$$\dot{x}(t) = f(t, x(t)), \quad t \in [t_0, t_f],$$

$$\text{und} \quad x(t_0) = x_0$$

eine Lösung hat, und diese Lösung ist eindeutig, d.h. es gibt keine andere Kurve, die Lösung des AWP ist.

Der Satz ist in Abbildung 8.2 illustriert. Wir wissen also, dass es theoretisch eine Lösung des AWP gibt, wenn die Funktion f stetig differenzierbar ist (was sie in fast allen praktischen Fällen auch ist) und wir im Inneren ihres Definitionsbereiches starten. Aber wie löst man das AWP

Abbildung 8.2: Existenz und Eindeutigkeit bis an den Rand des Gebietes V

praktisch? Anschaulich ginge dies so, dass man die Kurve findet, die in x_0 startet und dann immer tangential zu $f(t, x)$ verläuft, dass man so etwas wie eine „Stromlinie“ des Vektorfeldes $f(t, x)$ berechnet.

Aber wie löst man dieses Problem mathematisch? Es ist im allgemeinen schwer, eine analytisch darstellbare Lösungskurve zu einem AWP anzugeben, aber man kann einige günstige Spezialfälle leicht behandeln. Für alle anderen Fälle kann man mit Hilfe des Computers näherungsweise Lösungen berechnen, mit Methoden, die die sogenannte *Numerik* bereitstellt. In diesem Skript wollen wir jedoch zunächst einige günstige Spezialfälle betrachten.

8.1 Systeme mit einer Zustandsvariablen

Falls $n = 1$, also $x(t) \in \mathbb{R}$, kann man sich alles leicht veranschaulichen. Die Zahl $f(t, x)$ entspricht der Steigung, die die Kurve $x(t)$, an der Stelle (t, x) haben soll, siehe Abbildung 8.3. Es gibt einen sehr einfach zu lösenden Spezialfall, den wir hier kurz behandeln wollen.

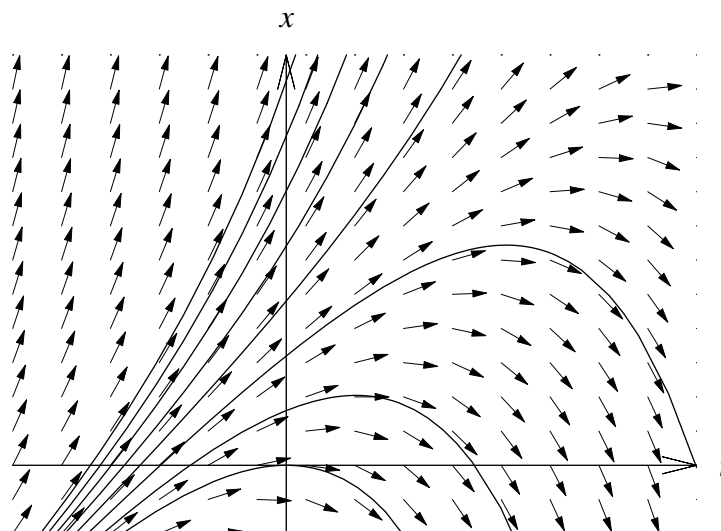


Abbildung 8.3: Lösungen der Gleichung $\dot{x} = x - t$ und die Steigung als Funktion von Zustand und Zeit.

Separation der Variablen

Falls $f(t, x) = g(x) \cdot h(t)$ formt man wie folgt um

$$\frac{dx}{dt} = g(x) \cdot h(t) \Leftrightarrow \frac{dx}{g(x)} = dt \cdot h(t) \quad (8.4)$$

$$\int_{x_0}^{x(t)} \frac{dx}{g(x)} = \int_{t_0}^t h(t) \cdot dt \quad (8.5)$$

Beispiel 8.1.1 (exponentielles Wachstum).

$$\begin{aligned} \dot{x} &= a \cdot x \\ \frac{dx}{dt} &= a \cdot x \Leftrightarrow \frac{dx}{x} = a \cdot dt \\ \int_{x_0}^{x(t)} \frac{dx}{x} &= a \int_{t_0}^t dt \Leftrightarrow \ln x(t) - \ln x_0 = a(t - t_0) \\ &\Leftrightarrow \ln \frac{x(t)}{x_0} = a(t - t_0) \\ &\Leftrightarrow x(t) = x_0 \cdot e^{a(t-t_0)} \end{aligned}$$

Man sieht, dass für $a > 0$ jede Lösungskurve exponentiell wächst, und für $a < 0$ exponentiell abfällt. Die Gleichung $\dot{x} = a \cdot x$ beschreibt z.B. das Wachstum von Bakterien, das Anwachsen festverzinsten Geldes, den Abbau von Medikamenten im Körper, radioaktiven Zerfall, usw. In

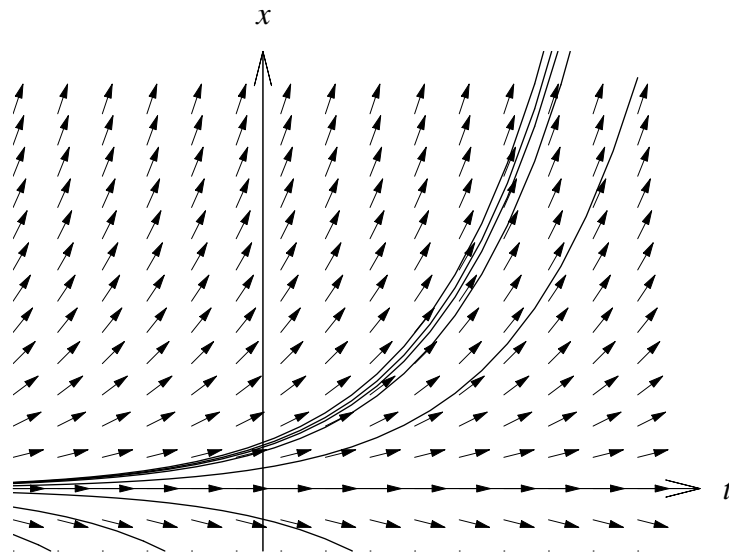


Abbildung 8.4: Lösungen der Gleichung $\dot{x} = x$. Man beachte, dass die Steigungspfeile nicht von t abhängen, im Gegensatz zu Abbildung 8.3.

Abbildung 8.4 sind für $a = 1$ die Steigung als Funktion des Ortes und einige Lösungskurven dargestellt.

8.2 Der harmonische Oszillator

Wir betrachten nun zweidimensionale Systeme:

$$\begin{aligned}\dot{x}_1(t) &= f_1(t, x_1(t), x_2(t)) \\ \dot{x}_2(t) &= f_2(t, x_1(t), x_2(t)).\end{aligned}$$

Diese Systeme kann man sich im autonomen Fall noch ganz gut veranschaulichen, indem man das Vektorfeld auf die Ebene einzeichnet, siehe Abbildung 8.6. Allgemeine Lösungsmethoden zur Lösung von AWP's gibt es jedoch nur für Spezialfälle, von denen wir nur zwei der allerwichtigsten, den *harmonischen* und den *gedämpften harmonischen Oszillator* in diesem Abschnitt behandeln wollen. Wir beginnen mit einem Beispiel zur Motivation.

Beispiel 8.2.1 (Federpendel).

Seien p die Position und v die Geschwindigkeit einer Masse an einer Feder, siehe Abbildung 8.5. Wenn die Masse m ausgelenkt wird, gibt es eine rücktreibende Kraft $F = -kp$, die eine Beschleunigung $\dot{v} = F/m$ bewirkt. Das System gehorcht den Differentialgleichungen:

$$\begin{aligned}\dot{p}(t) &= v(t) \\ \dot{v}(t) &= \frac{F}{m} = -\frac{k \cdot p(t)}{m} = -c \cdot p(t)\end{aligned}\tag{8.6}$$

wobei $c = \frac{k}{m} \quad \Leftarrow \text{Federkonstante}$
 $\quad \quad \quad \quad \quad \Leftarrow \text{Masse}$

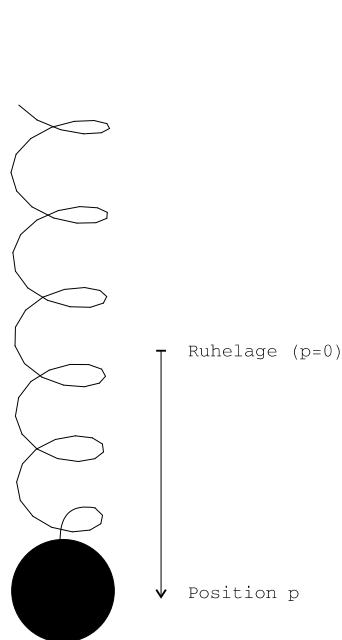


Abbildung 8.5: Federpendel: die Federkraft ist proportional zu p .

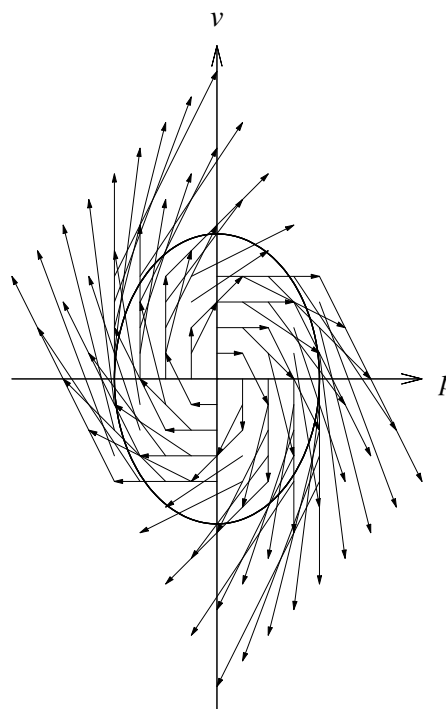


Abbildung 8.6: Vektorfeld des harmonischen Oszillators für $c = 2$ und eine ellipsenförmige Lösungskurve.

Im allgemeinen bezeichnet man als harmonischen Oszillator jedes autonome System, dass durch die folgenden Differentialgleichungen beschrieben wird:

$$\dot{x}_1 = x_2 \quad (8.7)$$

$$\dot{x}_2 = -c \cdot x_1 \quad (8.8)$$

$$\text{also } \dot{x} = f(x) \quad \text{mit } f(x) = \begin{pmatrix} x_2 \\ -cx_1 \end{pmatrix}.$$

Eine Veranschaulichung des Vektorfelds $f(x) = \begin{pmatrix} x_2 \\ -cx_1 \end{pmatrix}$ geben wir in Abbildung 8.6.

8.2.1 Lösungsansatz im Reellen

Wie erhalten wir nun Lösungskurven, die die Differentialgleichungen (8.7) und (8.8) erfüllen? Wie finden wir also die *Stromlinien* des Vektorfeldes $f(x)$, das in Abbildung 8.6 veranschaulicht ist? Falls wir $x_1(t)$ kennen, erhalten wir mit Gleichung (8.7) sofort auch $x_2(t) = \dot{x}_1(t)$. Aus Abbildung 8.6 erraten wir, dass die Lösungskurven Ellipsen sein könnten? Wir machen also einfach einmal den Ansatz:

$$x_1(t) = a \cdot \sin(\omega t).$$

Daraus folgt

$$\dot{x}_1(t) = a \cdot \omega \cdot \cos(\omega t) = x_2(t)$$

nach (8.7) und somit

$$\dot{x}_2(t) = -a \omega^2 \sin(\omega t).$$

Andererseits folgt mit (8.8):

$$\dot{x}_2(t) = -c x_1(t) = -a c \sin(\omega t).$$

Dies geht nur, wenn $\omega := \sqrt{c}$. Die *Amplitude* a kann beliebig sein. Ebenso gibt der Ansatz $x_1(t) = \tilde{a} \cdot \cos(\omega t)$ eine Lösung, mit beliebigem $\tilde{a}$. Wir werden in Abschnitt 8.3 rigoros zeigen, dass für *lineare* Systeme, wie es auch der harmonische Oszillator ist, die Linearkombination zweier Lösungen selbst wieder eine Lösung ist. Deshalb können wir den allgemeinen Ansatz

$$x_1(t) = a \sin(\omega t) + \tilde{a} \cos(\omega t)$$

machen. Dies erlaubt uns schliesslich, ein Anfangswertproblem mit einem beliebigen, aber festen Anfangswert $x_1(0), x_2(0)$ zu lösen. Wir verwenden unseren Ansatz und vergleichen:

$$\begin{aligned} x_1(0) &= a \cdot \sin(\omega t) + \tilde{a} \cos(\omega t)|_{t=0} = \tilde{a} \\ x_2(0) &= \omega(a \cdot \cos(\omega t) - \tilde{a} \sin(\omega t))|_{t=0} = \omega a \end{aligned}$$

Also ist mit

$$\tilde{a} = x_1(0) \quad \text{und} \quad a = \frac{x_2(0)}{\omega} \tag{8.9}$$

tatsächlich durch $x_1(t) = a \sin(\omega t) + \tilde{a} \cos(\omega t)$ eine Lösung des AWP gegeben. Zusammenfassend ergibt sich also: die allgemeine Lösung des AWP

$$\dot{x}_1 = x_2, \quad \dot{x}_2 = -c x_1, \quad x(0) = \begin{pmatrix} x_1(0) \\ x_2(0) \end{pmatrix}, \tag{8.10}$$

ist mit $\omega := \sqrt{c}$ durch

$$\begin{aligned} x_1(t) &= \frac{x_2(0)}{\omega} \cdot \sin(\omega t) + x_1(0) \cdot \cos(\omega t) \\ x_2(t) &= \dot{x}_1(t) = x_2(0) \cos(\omega t) - x_1(0) \cdot \omega \cdot \sin(\omega t) \end{aligned}$$

gegeben.

8.2.2 Lösungsansatz im Komplexen

Oft ist es praktisch, bei der Lösung von (linearen) gewöhnlichen Differentialgleichungen komplexe Zahlen zu verwenden. Dies ist zwar etwas abstrakter, aber oft einfacher, insbesondere beim gedämpften Oszillator, den wir im Abschnitt 8.2.3 behandeln werden. Wir stellen diesen Ansatz im Komplexen jetzt vor, indem wir damit noch einmal den harmonischen Oszillator behandeln. Wir machen den einfachen Ansatz

$$x_1(t) = e^{\lambda t}$$

und daraus folgt wieder

$$\begin{aligned}x_2(t) &= \dot{x}_1(t) = \lambda e^{\lambda t} \\ \dot{x}_2(t) &= \lambda^2 e^{\lambda t} = -c x_1(t) = -c e^{\lambda t} \\ \Leftrightarrow \lambda^2 &= -c \\ \Leftrightarrow \lambda &= \pm i\sqrt{c} = \pm i\omega\end{aligned}$$

wobei i die imaginäre Einheit ist. Daraus folgt, dass auch

$$\begin{pmatrix} x_1(t) \\ x_2(t) \end{pmatrix} = \begin{pmatrix} e^{i\omega t} \\ i\omega e^{i\omega t} \end{pmatrix} \quad \text{und} \quad \begin{pmatrix} e^{-i\omega t} \\ -i\omega e^{-i\omega t} \end{pmatrix}$$

Lösungskurven des gewöhnlichen Differentialgleichungssystems (8.7) und (8.8) sind. Wir werden sehen, dass diese Lösungen im Prinzip genau die gleichen sind, wie die, die wir zuvor im Reellen erhalten haben. Es lässt sich wieder durch Linearkombination mit (hier komplexen) Faktoren a_1, a_2 die Lösung des AWP (8.10) konstruieren

$$\begin{aligned}x_1(t) &= a_1 e^{i\omega t} + a_2 e^{-i\omega t} \Rightarrow x_1(0) = a_1 + a_2 \\ x_2(t) &= i\omega a_1 e^{i\omega t} - i\omega a_2 e^{-i\omega t} \Rightarrow x_2(0) = i\omega(a_1 - a_2)\end{aligned}$$

Auflösen ergibt, dass mit

$$a_1 = \frac{1}{2} \left(x_1(0) + \frac{x_2(0)}{i\omega} \right) \quad \text{und} \quad a_2 = \frac{1}{2} \left(x_1(0) - \frac{x_2(0)}{i\omega} \right) \quad (8.11)$$

die Lösung des AWP (8.10) auch durch $x_1(t) = a_1 e^{i\omega t} + a_2 e^{-i\omega t}$ dargestellt werden kann.

Vergleich mit der reellen Lösung: Der Ansatz im Komplexen ist analog zur reellen Lösung des AWP, nur dass statt $x_1(t) = \sin \omega t$ und $\cos \omega t$ jetzt $e^{i\omega t}$ und $e^{-i\omega t}$ ein Lösungspaar sind, mit dessen Hilfe wir das AWP lösen können. Die auf beide Weisen erhaltenen Lösungen sind tatsächlich identisch, denn mit den Koeffizienten $a, \tilde{a}$ aus (8.9) und denen aus (8.11) gilt

$$x_1(t) = a \sin \omega t + \tilde{a} \cos \omega t = a_1 e^{i\omega t} + a_2 e^{-i\omega t},$$

wie man leicht unter Verwendung der Euler-Formel

$$e^{i\alpha} = \cos \alpha + i \sin \alpha$$

nachprüft, die umgeformt

$$\cos \alpha = \frac{e^{i\alpha} + e^{-i\alpha}}{2} \quad \text{und} \quad \sin \alpha = \frac{e^{i\alpha} - e^{-i\alpha}}{2i}$$

ergibt.

8.2.3 Der gedämpfte harmonische Oszillator

In fast allen real vorkommenden Oszillatoren gibt es Energieverluste, die dazu führen, dass die Amplitude mit der Zeit abklingt; man spricht dann von *Dämpfung*. Im Federpendel aus Beispiel 8.2.1 wird die Bewegung beispielsweise gedämpft, weil es Reibungsverluste gibt. Statt der Gleichung (8.6) $\dot{v}(t) = -\frac{k}{m}p(t)$ gilt nun

$$\dot{v}(t) = -\frac{k \cdot p(t)}{m} - \beta v(t),$$

wobei der Reibungsterm $-\beta v(t)$ (mit einer Konstanten β) zu einer der Geschwindigkeit proportionalen Abbremsung führt.

Allgemein nennt man jedes System, das durch die Gleichungen

$$\dot{x}_1 = x_2 \tag{8.12}$$

$$\dot{x}_2 = -cx_1 - \beta x_2 \tag{8.13}$$

beschrieben wird, einen **gedämpften harmonischen Oszillator**. Eine Veranschaulichung des Vektorfelds $f(x) = \begin{pmatrix} x_2 \\ -cx_1 - \beta x_2 \end{pmatrix}$ ist in Abbildung 8.7 zu sehen.

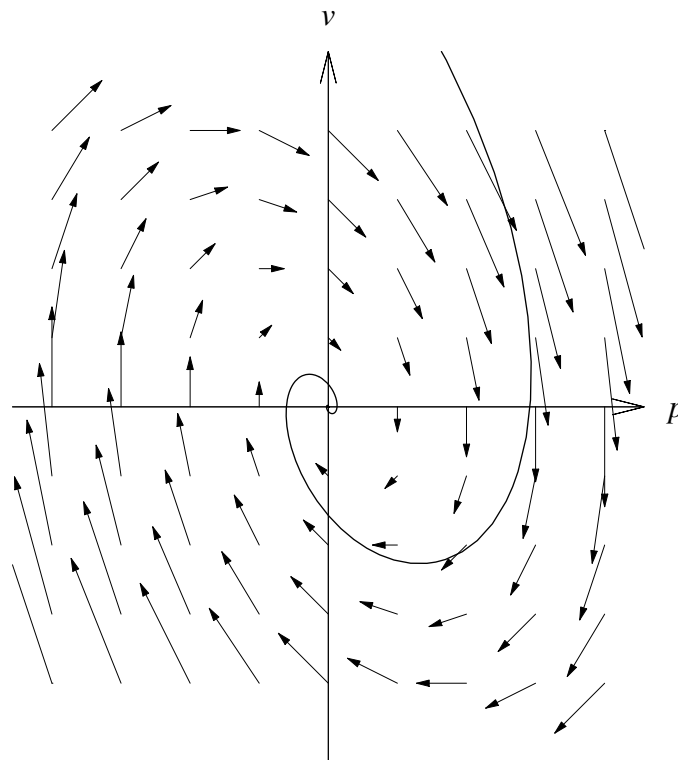


Abbildung 8.7: Das Vektorfeld des gedämpften harmonischen Oszillators und eine spiralförmige Lösungskurve.

8.2.4 Lösungsansatz im Komplexen

Wir verwenden wieder den Ansatz im Komplexen, $x_1(t) = e^{\lambda t}$, der sich für den gedämpften harmonischen Oszillator als wesentlich einfacher und eleganter herausstellt als die Rechnung im Reellen:

$$\begin{aligned} x_1(t) = e^{\lambda t} &\Rightarrow x_2 = \lambda e^{\lambda t} \Rightarrow \dot{x}_2 = \lambda^2 e^{\lambda t} \stackrel{!}{=} -c e^{\lambda t} - \beta \lambda e^{\lambda t} \\ &\Leftrightarrow \lambda^2 = -c - \beta \lambda \\ &\Leftrightarrow \lambda^2 + \beta \lambda + c = 0, \end{aligned}$$

und diese Gleichung hat die beiden Nullstellen

$$\lambda_{1,2} = -\frac{\beta}{2} \pm \sqrt{\frac{\beta^2}{4} - c}.$$

Wie die entsprechenden Lösungen $e^{\lambda t}$ aussehen, hängt nun entscheidend vom Vorzeichen des Terms $\frac{\beta^2}{4} - c$ unter der Wurzel ab.

Fall 1: $\beta^2 < 4c$ (schwache Dämpfung)

$$\pm \sqrt{\frac{\beta^2}{4} - c} = \pm i \sqrt{c - \frac{\beta^2}{4}} =: \pm i\omega$$

D.h.

$$\begin{aligned} x_1(t) &= e^{-\frac{\beta}{2}t} \cdot e^{+i\omega t} \text{ und } e^{-\frac{\beta}{2}t} \cdot e^{-i\omega t} \text{ sind Lösungen der Differentialgleichung, mit jeweils} \\ x_2(t) &= \dot{x}_1(t) = \left(-\frac{\beta}{2} \pm i\omega\right) \cdot e^{-\frac{\beta}{2}t} e^{\pm i\omega t} \end{aligned}$$

Im Reellen sind analog

$$x_1(t) = e^{-\frac{\beta}{2}t} \cdot \cos \omega t \quad \text{und} \quad = e^{-\frac{\beta}{2}t} \cdot \sin \omega t$$

Lösungen des Systems (8.12) und (8.13), siehe Abbildung 8.8. Zusammen mit dem zugehörigen $x_2(t) = \dot{x}_1(t)$ geben diese Lösungskurven tatsächlich Spiralen im (x_1, x_2) -Raum, wie in Abbildung 8.7 skizziert.

Fall 2: $\beta^2 > 4c$ (starke Dämpfung) $\sqrt{\frac{\beta^2}{4} - c}$ ist reell und wegen $\sqrt{\frac{\beta^2}{4} - c} \leq \sqrt{\frac{\beta^2}{4}} = \frac{\beta}{2}$ gilt

$$\lambda_{1,2} = -\frac{\beta}{2} \pm \sqrt{\frac{\beta^2}{4} - c} \leq 0 \quad \text{aber auch} \quad \lambda_{1,2} \geq -\beta.$$

Die allgemeine Lösung ist also

$$\begin{aligned} x_1(t) &= a_1 e^{\lambda_1 t} + a_2 e^{\lambda_2 t}, \quad \text{mit} \quad -\beta \leq \lambda_2 < \lambda_1 \leq 0, \quad \text{und} \\ x_2(t) &= \lambda_1 a_1 e^{\lambda_1 t} + \lambda_2 a_2 e^{\lambda_2 t}. \end{aligned}$$

Dies bedeutet einen exponentiellen Abfall mit zwei verschiedenen Zeitkonstanten, wie in Abbildung 8.9 skizziert.

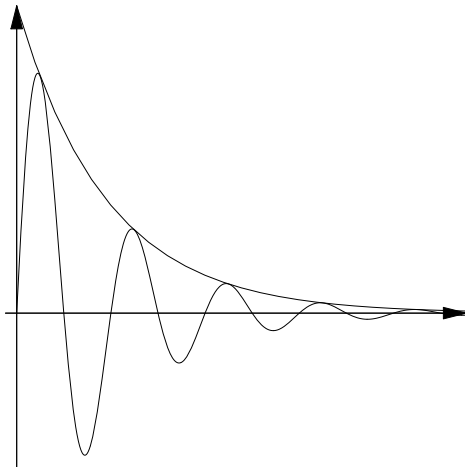


Abbildung 8.8: Lösung $x_1(t)$ des schwach gedämpften harmonischen Oszillators, über der Zeit aufgetragen (mit Einhüllender).

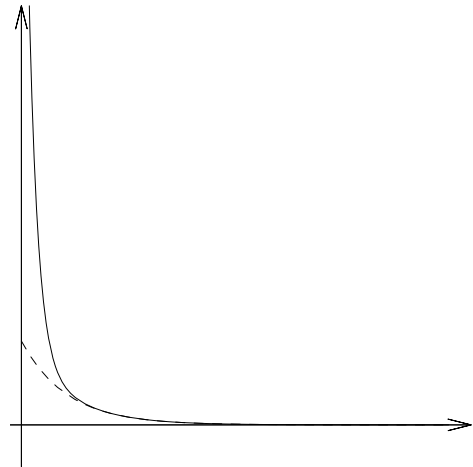


Abbildung 8.9: Exponentieller Abfall mit verschiedenen Zeitkonstanten beim stark gedämpften harmonischen Oszillator.

Lösung des allgemeinen Anfangswertproblems

Um beim gedämpften harmonischen Oszillator das Anfangswertproblem zu lösen, d.h. die Koeffizienten a_1 und a_2 für die Lösung $x_1(t) = a_1 e^{\lambda_1 t} + a_2 e^{\lambda_2 t}$ zu bestimmen, löst man ganz einfach wieder die Gleichungen

$$x_1(0) = a_1 + a_2 \quad (8.14)$$

$$x_2(0) = \lambda_1 a_1 + \lambda_2 a_2. \quad (8.15)$$

Interessanterweise gilt dieser Ansatz auch für den ungedämpften harmonischen Oszillator, bei dem einfach $\lambda_1 = i\omega$ und $\lambda_2 = -i\omega$ ist. Wir werden gleich sehen, dass die Zeitkonstanten $\lambda_{1,2}$ des komplexen Ansatzes auch als Eigenwerte der sogenannten **Systemmatrix** A eines linearen Systems aufgefasst werden können.

8.3 Lineare dynamische Systeme

Wir liefern nun nachträglich etwas Theorie, die uns erlaubt, nicht nur den (gedämpften) harmonischen Oszillator besser zu verstehen, sondern für die äußerst wichtige Klasse der *linearen dynamischen Systeme* ein ganz allgemeines Lösungsverfahren anzugeben.

Definition 8.3.1 (Lineares Dynamisches System).

Falls

$$f(t, x) = A(t) \cdot x$$

mit $A(t) \in \mathbb{R}^{n \times n}$ sagt man, das System $\dot{x} = f(t, x)$ ist **linear**.

Es gilt nun

Satz 8.3.2 (Linearkombination von Lösungen).

Falls $y(t), z(t)$ Lösungen einer linearen Systemgleichung $\dot{y} = f(t, y)$ sind, dann ist auch jede Linearkombination $x(t) := \lambda_1 y(t) + \lambda_2 z(t)$ Lösung der Systemgleichung.

Beweis:

$$\begin{aligned}\dot{x}(t) &= \lambda_1 \dot{y}(t) + \lambda_2 \dot{z}(t) = \lambda_1 f(t, y(t)) + \lambda_2 f(t, z(t)) \\ &= \lambda_1 A(t)y(t) + \lambda_2 A(t)z(t) \\ &= A(t)(\lambda_1 \cdot y(t) + \lambda_2 \cdot z(t)) = A(t) \cdot x(t) \\ &= f(t, x(t))\end{aligned}\quad \square$$

Dies ist die nachträgliche Berechtigung für unsere Methode, das AWP des harmonischen Oszillators durch Linearkombination zweier Lösungen zu behandeln.

Uns interessiert meist nur der autonome Fall $\dot{x} = A \cdot x$ mit konstanter Matrix $A \in \mathbb{R}^{n \times n}$, für den wir bereits zwei Beispiele kennen:

1. $\dot{x} = a \cdot x$ (Zerfall oder Wachstum)
2. $\begin{pmatrix} \dot{x}_1 \\ \dot{x}_2 \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ -c & -\beta \end{pmatrix} \cdot \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$ (gedämpfter Oszillator, bzw. mit $\beta = 0$ ungedämpft)

Für autonome lineare Systeme gilt nun die folgende sehr mächtige Aussage:

Lemma 8.3.3 (Eigenwerte als Zeitkonstanten).

Sei $\dot{x} = A \cdot x$ ein autonomes lineares System und v sei Eigenvektor der Matrix A mit Eigenwert λ , d.h. $A \cdot v = \lambda \cdot v$. Dann ist $x(t) = e^{\lambda \cdot t} \cdot v$ eine Lösungskurve des linearen Differentialgleichungssystems.

Beweis: $\dot{x}(t) = \lambda \cdot e^{\lambda t} \cdot v = A \cdot e^{\lambda t} \cdot v = A \cdot x$. $\square$

Beispiel 8.3.4. Wir betrachten $\dot{x} = A \cdot x$ mit

$$A = \begin{pmatrix} a & 0 \\ 0 & b \end{pmatrix}, \quad v = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad \text{also} \quad A \cdot v = a \cdot v.$$

Der Satz besagt nun, dass

$$x(t) = e^{at} \cdot v = e^{at} \cdot \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \begin{pmatrix} e^{at} \\ 0 \end{pmatrix}$$

eine Lösung der linearen gewöhnlichen Differentialgleichung ist. Tatsächlich erhalten wir durch Einsetzen von $x(t)$ in die zwei Systemgleichungen

$$\begin{aligned}\dot{x}_1(t) &= a \cdot e^{at} = a \cdot x_1(t) \\ \dot{x}_2(t) &= 0 = b \cdot x_2(t).\end{aligned}$$

Beispiel 8.3.5. Als ein zweites Beispiel betrachten wir einen stark gedämpften Oszillator, mit $c = 1, \beta = 2$.

$$A = \begin{pmatrix} 0 & 1 \\ -1 & -2 \end{pmatrix}, \quad A \cdot \begin{pmatrix} 1 \\ -1 \end{pmatrix} = \begin{pmatrix} -1 \\ 1 \end{pmatrix} = \underbrace{(-1)}_{=\lambda} \cdot \underbrace{\begin{pmatrix} 1 \\ -1 \end{pmatrix}}_{=v}.$$

Der Satz besagt nun, dass

$$x(t) = e^{-t} \cdot \begin{pmatrix} 1 \\ -1 \end{pmatrix} = \begin{pmatrix} e^{-t} \\ -e^{-t} \end{pmatrix}$$

Lösung der Systemgleichungen $\dot{x} = A \cdot x$ ist. Wir testen dies durch Einsetzen

$$\begin{aligned} \dot{x}_1(t) &= -e^{-t} = x_2(t), \\ \dot{x}_2(t) &= e^{-t} = -e^{-t} + 2e^{-t} = -1 \cdot x_1(t) - 2x_2(t). \end{aligned}$$

Bemerkung 8.3.6 (Bedeutung komplexer Eigenwerte).

Falls λ nicht reell ist, also $\lambda = \alpha + i\omega$, dann gibt es je nach Vorzeichen von α eine auf- oder abschwellige Oszillation, denn

$$e^{\lambda t} = \underbrace{e^{\alpha t}}_{\text{Dämpfung oder Wachstum}} \cdot \underbrace{e^{i\omega t}}_{\text{Oszillation}}.$$

Wenn λ rein imaginär ist, also $\alpha = 0$, hat die Oszillation eine konstante Amplitude.

Beispiel 8.3.7 (Schwach gedämpfter Oszillator). $\beta^2 < 4c$

Wir betrachten die Systemmatrix

$$A = \begin{pmatrix} 0 & 1 \\ -c & -\beta \end{pmatrix}.$$

$$\begin{aligned} \lambda \text{ Eigenwert von } A &\Leftrightarrow \det(A - \lambda \mathbb{I}) = 0 \\ &\Leftrightarrow \det \begin{pmatrix} -\lambda & 1 \\ -c & -\beta - \lambda \end{pmatrix} = 0 \\ &\Leftrightarrow \lambda^2 + \beta\lambda + c = 0 \\ &\Leftrightarrow \lambda_{1,2} = -\frac{\beta}{2} \pm \sqrt{\frac{\beta^2}{4} - c} = -\frac{\beta}{2} \pm i\sqrt{c - \frac{\beta^2}{4}}. \end{aligned}$$

Die Gleichung und die beiden Größen $\lambda_{1,2}$ kennen wir bereits! Wir haben es also mit zwei Eigenwerten $\lambda_{1,2}$ zu tun, und beide haben einen negativen Realteil $\alpha < 0$, der ein Abklingen bedeutet, und einen nichtverschwindenden Imaginärteil $\omega \neq 0$, was bedeutet, dass das System oszilliert.

Wir können uns nun fragen, ob man, wenn man nur die Eigenwerte einer Systemmatrix A kennt, ganz allgemein die Lösung konstruieren kann. Dies geht tatsächlich, wie wir sogleich für den meist auftretenden Fall einer diagonalisierbaren Matrix A beweisen wollen.

Satz 8.3.8 (Allgemeine Lösung des AWP für lineare autonome Systeme).

Falls $A \in \mathbb{R}^{n \times n}$ diagonalisierbar ist, also

$$A = B \cdot D \cdot B^{-1}, \quad \text{mit} \quad D = \begin{pmatrix} \lambda_1 & & \\ & \ddots & \\ & & \lambda_n \end{pmatrix} \quad \text{und} \quad B = \left(v_1 \mid v_2 \mid \cdots \mid v_n \right),$$

kann *jede* Lösung von $\dot{x} = Ax$ durch

$$x(t) = \sum_{i=1}^n v_i \cdot e^{\lambda_i \cdot t} \cdot a_i$$

mit beliebigem Gewichtsvektor

$$a = \begin{pmatrix} a_1 \\ \vdots \\ a_n \end{pmatrix} \in \mathbb{R}^n$$

dargestellt werden. Insbesondere gilt

$$x(0) = \sum_{i=1}^n v_i \cdot a_i = B \cdot a,$$

so dass die Lösung des AWP

$$\dot{x}(t) = Ax(t), \quad \text{mit} \quad x(0) = x_0$$

durch die Gewichte $a = B^{-1} \cdot x_0 \in \mathbb{R}^n$ gegeben ist.

Wir geben für diesen Satz zwei Beispiele, zum ersten ein weiteres Mal den harmonischen Oszillator, und zum zweiten ein Modell aus der Pharmakokinetik.

Beispiel 8.3.9 (harmonischer Oszillator).

Man kann Satz 8.3.8 leicht auf den (gedämpften) harmonischen Oszillator anwenden, wo gilt:

$$A = \begin{pmatrix} 0 & 1 \\ -c & -\beta \end{pmatrix} = BDB^{-1} \quad \text{mit} \quad D = \begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix} \quad \text{sowie} \quad B = (v_1 | v_2) = \begin{pmatrix} 1 & 1 \\ \lambda_1 & \lambda_2 \end{pmatrix},$$

wie man leicht durch Bilden der Matrixprodukte Av_1, Av_2 unter Verwendung von $\lambda_i^2 = -c - \beta\lambda_i$ nachprüfen kann. Die Gewichte a_1, a_2 sind wie zuvor durch die Gleichungen (8.14)-(8.15) bestimmt, die wir jetzt schreiben als

$$x_0 = Ba \quad \Leftrightarrow \quad \begin{pmatrix} x_1(0) \\ x_2(0) \end{pmatrix} = \begin{pmatrix} 1 & 1 \\ \lambda_1 & \lambda_2 \end{pmatrix} \begin{pmatrix} a_1 \\ a_2 \end{pmatrix}.$$

Beispiel 8.3.10 (Pharmakokinetik: Medikamentenabbau im Körper).

Wir betrachten ein einfaches Modell, das beschreiben soll, wie ein Medikament, das sich in der Blutbahn befindet, durch die Niere abgebaut wird. Es besteht aus zwei Zustandsvariablen.

$$\begin{aligned} K(t) &= \text{Medikament im Körper} \\ U(t) &= \text{Medikament im Urin} \\ \dot{K}(t) &= -k \cdot K(t) \text{ (Ausscheidung durch die Niere)} \\ \dot{U}(t) &= +k \cdot U(t) \text{ (sammelt sich in der Blase)} \end{aligned}$$

Mit $x(t) = \begin{pmatrix} K(t) \\ U(t) \end{pmatrix}$ und $A = \begin{pmatrix} -k & 0 \\ k & 0 \end{pmatrix}$ erhalten wir $\dot{x}(t) = A \cdot x(t)$. Es gilt

$$A \cdot \begin{pmatrix} 1 \\ -1 \end{pmatrix} = \begin{pmatrix} -k \\ k \end{pmatrix} = -k \begin{pmatrix} 1 \\ -1 \end{pmatrix}, \quad \text{d.h. } \lambda_1 = -k \quad \text{und} \quad v_1 = \begin{pmatrix} 1 \\ -1 \end{pmatrix}, \quad \text{sowie}$$

$$A \cdot \begin{pmatrix} 0 \\ 1 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} = 0 \cdot \begin{pmatrix} 0 \\ 1 \end{pmatrix}, \quad \text{d.h. } \lambda_2 = 0 \quad \text{und} \quad v_2 = \begin{pmatrix} 0 \\ 1 \end{pmatrix}.$$

Also gilt

$$A = BDB^{-1} \quad \text{mit} \quad B = \begin{pmatrix} 1 & 0 \\ -1 & 1 \end{pmatrix} \quad \text{und} \quad D = \begin{pmatrix} -k & 0 \\ 0 & 0 \end{pmatrix},$$

denn $A \cdot B = B \cdot D \Leftrightarrow A = B^{-1} \cdot D \cdot B$. Man berechnet: zudem

$$B^{-1} = \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix}.$$

Für einen beliebigen Anfangswert K_0, U_0 ergibt sich also

$$a = B^{-1} \cdot x(0) = \begin{pmatrix} K_0 \\ K_0 + U_0 \end{pmatrix}.$$

Damit ergibt sich als Lösung des AWP:

$$\begin{aligned} x(t) &= \sum_{i=1}^2 v_i \cdot e^{\lambda_i t} a_i \\ &= \begin{pmatrix} 1 \\ -1 \end{pmatrix} e^{-kt} \cdot K_0 + \begin{pmatrix} 0 \\ 1 \end{pmatrix} \cdot e^0 \cdot (K_0 + U_0) \\ &= \begin{pmatrix} K_0 e^{-kt} \\ -K_0 e^{-kt} + (K_0 + U_0) \end{pmatrix}. \end{aligned}$$

Man erhält also

$$K(t) = K_0 \cdot e^{-kt} \quad \text{und} \quad U(t) = U_0 + K_0(1 - e^{-kt}),$$

wie in Abbildung 8.10 skizziert.

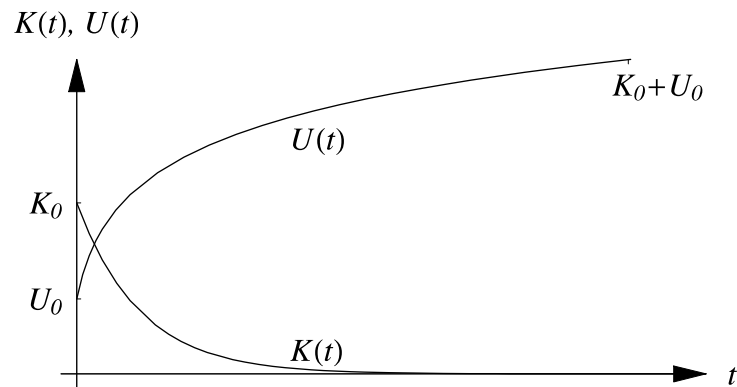


Abbildung 8.10: Ausscheidung eines Medikamentes in den Urin.

8.3.1 Stabilität und Eigenwerte

Wir wollen noch einen wichtigen Begriff kennenlernen, den der Stabilität.

Definition 8.3.11 (Stabilität, asymptotische Stabilität eines linearen Systems).

Ein lineares autonomes System $\dot{x} = Ax$ heißt **stabil**, falls es ein $C > 0$ gibt so dass für *alle* Lösungskurven $x(\cdot)$ gilt, dass

$$\sup_{t \rightarrow \infty} \|x(t)\| \leq C \|x(0)\|.$$

Ein lineares System heißt **asymptotisch stabil**, wenn für jede Trajektorie $x(\cdot)$ gilt

$$\lim_{t \rightarrow \infty} x(t) = 0,$$

ganz unabhängig vom Anfangswert $x(0)$.

Als ein Beispiel können wir uns den gedämpften harmonischen Oszillator vorstellen, dessen Trajektorien wegen des exponentiellen Dämpfungsterms $e^{-\frac{\beta}{2}t}$ alle gegen null konvergieren: er ist asymptotisch stabil. Der ungedämpfte harmonische Oszillator hingegen ist zwar stabil, denn die Trajektorien wachsen nicht über alle Grenzen, aber nicht asymptotisch stabil.

Die Stabilität hängt interessanterweise direkt mit den Eigenwerten der Systemmatrix A zusammen, wie wir für diagonalisierbare Matrizen A direkt aus Satz 8.3.8 folgern können:

Satz 8.3.12 (asymptotische Stabilität eines linearen Systems).

Ein lineares autonomes System $\dot{x} = Ax$ ist genau dann asymptotisch stabil, wenn alle Eigenwerte λ_i von A einen **negativen Realteil** haben.

Da die Eigenwerte eine so wichtige Rolle zur Charakterisierung des Systemverhaltens linearer autonomer Systeme haben, malt man sich oft, um eine Übersicht zu bekommen, die Eigenwerte in die komplexe Ebene, wie in Abbildung 8.11 für ein Beispielsystem mit 4 Eigenwerten dargestellt. Aus so einer Darstellung kann man einiges sehen:

- Sind alle Eigenwerte λ_i in der linken Halbebene, also mit negativem Realteil $\text{Re}(\lambda_i) < 0$, dann ist das System asymptotisch stabil.

- Außerdem gilt natürlich, dass jeder Eigenwert λ_i mit **nichtverschwindendem Imaginärteil** $\text{Im}(\lambda_i)$ eine Oszillation des Systems bedeutet, die man in der Praxis oft auch als **Resonanzfrequenz** bezeichnet. Resonanz tritt dann auf, wenn der entsprechende Eigenwert keinen zu stark negativen Realteil $\text{Re}(\lambda_i)$ hat, die Schwingung also nicht zu stark gedämpft ist, und wenn das System mit der Frequenz $\omega_i = \text{Im}(\lambda_i)$ angeregt wird. (Anregung dynamischer Systeme haben wir hier nicht behandelt, es geht im Wesentlichen um eine Änderung der Systemgleichungen zu $\dot{x} = Ax + \delta(t)$ mit einer periodischen Störung $\delta(t)$, z.B. dem wiederholten Anschubsen eines Federpendels.)
- Es ist interessant zu beobachten, dass für reelle Systemmatrizen $A \in \mathbb{R}^{n \times n}$ die komplexen Eigenwerte immer in konjugiert komplexen Paaren auftreten, also einmal über, einmal unter der reellen Achse im gleichen Abstand.
- Wenn ein lineares System lange Zeit ungestört bleibt, setzen sich die Komponenten, die zum Eigenwert (oder den Eigenwerten) mit dem größten Realteil gehören, durch, denn nach Satz 8.3.8 gilt, falls $\text{Re}(\lambda_1) > \text{Re}(\lambda_i)$ für $i = 2, \dots, n$:

$$x(t) = e^{\lambda_1 t} \cdot \left(v_1 a_1 + \underbrace{\sum_{i=2}^n e^{(\lambda_i - \lambda_1)t} v_i a_i}_{\rightarrow 0 \text{ für } t \rightarrow \infty} \right) \approx \underbrace{e^{\lambda_1 t} \cdot v_1 \cdot a_1}_{\text{„dominante Komponente“}} \quad (8.16)$$

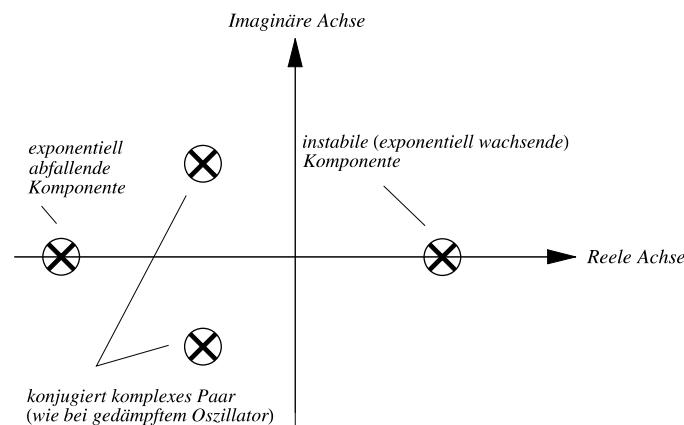


Abbildung 8.11: Die Eigenwerte eines linearen autonomen Systems in der komplexen Ebene.

Bemerkung 8.3.13. Wir erwähnen hier, dass auf der Theorie linearer Systeme ein ganzes Spezialgebiet der Ingenieurwissenschaften aufbaut, die *klassische Regelungstheorie*, die äußerst stark von Techniken der linearen Algebra wie z.B. der hier vorgestellten vorgestellten Eigenwertanalyse und von komplexen Zahlen Gebrauch macht. Die Eigenwerte von A heißen bei den Regelungstechnikern meist die „Pole“ des Systems. Typischerweise werden dort Systeme der Form $\dot{x} = Ax + Bu(t)$ betrachtet, mit dem System von außen vorgegebenen *Kontrollen* $u(t)$.

8.4 Nichtlineare autonome Systeme

In den meisten Anwendungen ist die Funktion $f(x)$ in dem System von Differentialgleichungen $\dot{x} = f(x)$ **nichtlinear** !

In diesen Fällen gibt es kein allgemeines Lösungsverfahren und man ist in der Regel auf numerische Approximation von Lösungen mit Hilfe eines Computers angewiesen.

Im Rahmen einer “qualitativen” Theorie interessiert man sich häufig für spezielle Lösungen. Dies sind meist “konstante” Lösungen (sog. Fixpunkte mit $f(x) = 0$) oder oszillierende Lösungen (sog. periodische Orbits).

Definition 8.4.1 (periodischer Orbit). Eine Lösung $x(t)$ des dynamischen Systems $\dot{x} = f(x)$ heißt periodisch, falls es ein $T \in \mathbb{R}, T < \infty$ gibt, so dass $x(t + T) = x(t) \quad \forall t \in \mathbb{R}, t > 0$.

In nichtlinearen Systemen ist komplexes dynamisches Verhalten von Lösungen möglich: z.B. permanent oszillierende Lösungen oder sog. deterministisches Chaos (s.u.).

Allgemein ist es schwierig, *nichtlineare* Systeme zu analysieren. Ihr Verhalten kann beliebig komplex werden, und sogar zu *deterministischem Chaos* führen, das ist ein Systemverhalten, das bei kleinen Änderungen des Anfangswertes x_0 nach einiger Zeit gänzlich verschiedene Lösungskurven erzeugt. Wir geben für Interessierte dafür ein Beispiel mit einer kleinen Anleitung, wie man ganz allgemeine nichtlineare Anfangswertprobleme der Form

$$\dot{x}(t) = f(t, x(t)), \quad \text{mit} \quad x(t_0) = x_0$$

mit Hilfe des Computers lösen kann.

***Beispiel 8.4.2** (Lorenz-Attraktor).

Im Jahre 1963 fand der Meteorologe Ed N. Lorenz ein relativ einfaches System von 3 gewöhnlichen Differentialgleichungen, mit dem er ursprünglich versucht hatte, die Konvektion in der Erdatmosphäre zu modellieren, das aber ein äußerst seltsames, „chaotisches“, Verhalten zeigte. Dieses System, das heute auch der **Lorenz-Attraktor** genannt wird, ist beschrieben durch:

$$\begin{aligned}\dot{x}_1 &= a(x_2 - x_1) \\ \dot{x}_2 &= x_1(b - x_3) - x_2 \\ \dot{x}_3 &= x_1x_2 - cx_3\end{aligned}$$

wobei $a = 10, b = 28, c = \frac{8}{3}$.

Man kann mit Hilfe von SCILAB (bzw. MATLAB) das System für einen gegebenen Anfangswert simulieren, z.B. für $x_0 = (1, 1, 1)^T$, indem man die Anfangswertproblem-Lösungsroutine ode (bzw. ode45) verwendet. Dafür müssten Sie zunächst die Systemfunktion definieren, so dass die Gleichungen die Form $\dot{x} = f(t, x)$ haben.

Wir nennen diese Funktion `f_lorenz` und schreiben

```
function [xdot]=f_lorenz(t,x)
a=10; b=28; c=8/3;
xdot=zeros(3,1);
```

```

xdot(1) = a * (x(2) - x(1));
xdot(2) = x(1) * (b - x(3)) - x(2);
xdot(3) = x(1) * x(2) - c * x(3);
endfunction

```

Nun laden wir die Funktion mit `getf` und rufen den AWP Löser mit den Zeilen

```

x0=[1 1 1]';
xmat=ode(x0, 0, [0:0.01:50], f_lorenz);

```

auf. In MATLAB müßte man die Funktion unter dem Namen `f_lorenz.m` abspeichern und tippen:

```

x0=[1 1 1]';
[t,xmat]=ode45('f_lorenz',[0:0.01:50],x0'); xmat=xmat';

```

Dies liefert uns die Lösungskurve, die bei x_0 zur Zeit $t = 0$ startet, an den Stellen $t = 0.00, 0.01, 0.02, \dots, 50.00$, als eine 3×5001 -Matrix. Mit `plot([0:0.01:50],xmat(1,:))` können Sie sich die Werte für $x_1(t)$ gegen die Zeit ansehen, und mit `plot(xmat(1,:),xmat(2,:))` können Sie sich die Figur in der x_1, x_2 -Ebene ansehen. Wir machen nun das Experiment, die Routine für den Anfangswert $x_0 = (1.001, 1, 1)^T$ aufzurufen. Die beiden Ergebnisse sind in Abbildung 8.12 gezeigt.

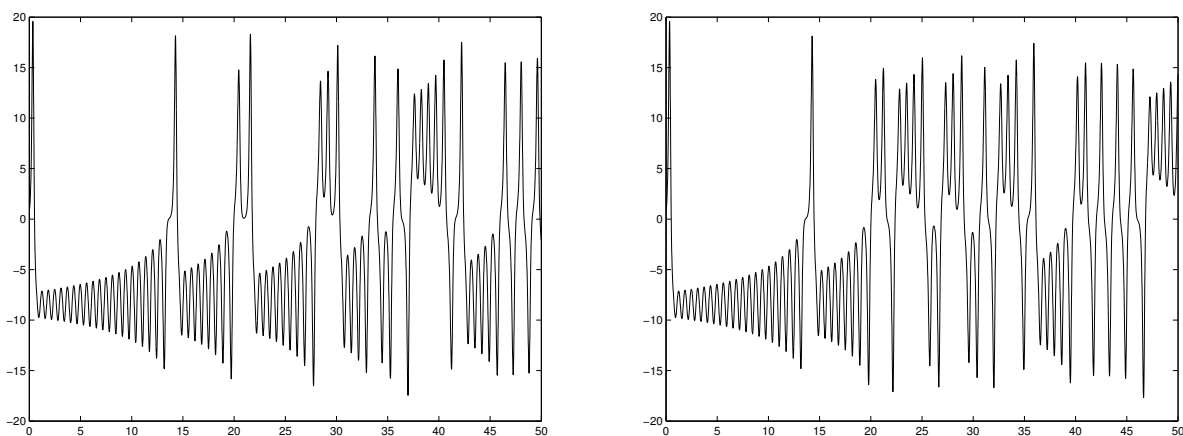


Abbildung 8.12: Lösungskurven $x_1(t)$ des Lorenz-Attraktors, für die Anfangswerte $x_0 = (1, 1, 1)^T$ und $x_0 = (1.001, 1, 1)^T$. Nach etwa 20 Zeiteinheiten werden sie sehr verschieden.

8.4.1 Fixpunkte und Stabilität

Da allgemeine Aussagen über die Lösungen nichtlinearer dynamischer Systeme schwer zu erhalten sind, erwähnen wir hier nur eine sehr wichtige Technik, die sich Techniken der linearen Systemtheorie zunutze macht; sie hilft, das Verhalten in der Nähe sogenannter *Fixpunkte* zu verstehen.

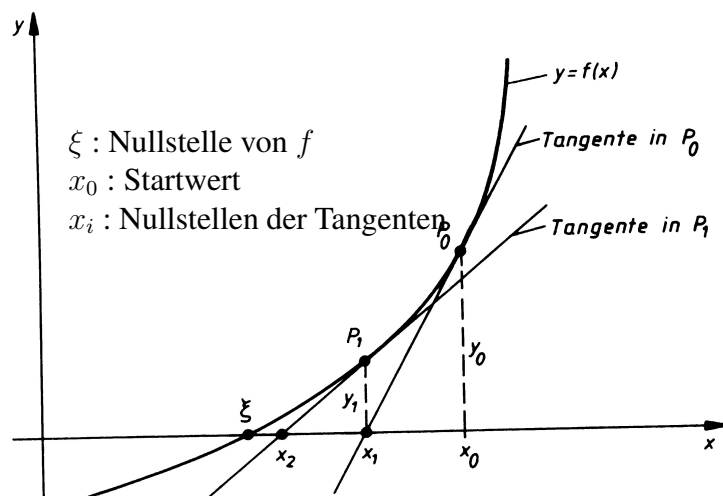
Definition 8.4.3 (Fixpunkt).

Ein Vektor $x^* \in \mathbb{R}^n$ heisst **Fixpunkt** des dynamischen Systems $\dot{x} = f(x)$, wenn $f(x^*) = 0$.

Anschaulich bedeutet dies: wenn man mit $x(0) = x^*$ startet, bleibt die Trajektorie für immer im Fixpunkt, also $x(t) = x^*$. Die Berechnung von Fixpunkten ist für nichtlineare Systeme nur selten analytisch möglich, daher kommen numerische Verfahren zum Einsatz.

Das Newton-Verfahren

Ziel des Newton-Verfahrens ist die Berechnung von Nullstellen nichtlinearer Gleichungssysteme $f(x) = 0$. Es wird daher auch zur Berechnung von Fixpunkten nichtlinearer Differentialgleichungssysteme eingesetzt. Motivation des Newton-Verfahrens für eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$:



→ Bestimme sukzessive Nullstellen von Tangenten an den Funktionsgraphen von f .

Mathematische Formulierung für $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$:

Wähle Startwert x^0 (Schätzung für Nullstelle), (x^0 "oberer" Index für Folge im $\mathbb{R}^n$)

Taylor-Entwicklung 1.Ordnung von f um x^0 :

$$f(x) \approx f(x^0) + \underbrace{f'(x^0)}_{\text{Jacobi-Matrix}} (x - x^0) \quad (\text{Gleichung der "Tangente" in } x)$$

Suche Nullstelle x der "Tangente":

$$\begin{aligned}
 f(x^0) + f'(x^0)(x - x^0) &= 0 \\
 \Leftrightarrow x &= x^0 - \underbrace{(f'(x^0))^{-1} f(x^0)}_{\text{Inverse der Jacobi-Matrix}}, \quad \text{setze } x^1 := x
 \end{aligned}$$

und iteriere das Verfahren (Taylor-Entwicklung um x^1 , Bestimmung Nullstelle der "Tangenten" in $x^1 \rightarrow x^2, \dots$)

In der Praxis berechnet man die Inverse der Jacobi-Matrix nicht explizit, sondern löst das Lineare Gleichungssystem

$$f'(x^0)\Delta x^0 = -f(x^0), \Delta x^0 = x - x^0$$

und setzt $x^1 := x^0 + \Delta x^0, \dots,$

$$x^n := x^{n-1} + \Delta x^{n-1}, f'(x^{n-1})\Delta x^{n-1} = -f(x^{n-1})$$

Die Iteration (Folge) $x^n, n \rightarrow \infty$ heisst konvergent gegen die Nullstelle $\bar{x}$ von f , falls gilt: $\lim_{n \rightarrow \infty} x^n = \bar{x}$

Satz 8.4.4 (Konvergenzsatz). Die Newton-Iteration im $\mathbb{R}$ konvergiert, falls im Intervall $[a, b]$, in dem alle Iterierten x^i ($i = 0, \dots, n$) liegen, gilt:

$$\frac{f(x)f''(x)}{[f'(x)]^2} < 1(*)$$

(Spezialfall 1-dem. Problem $f : \mathbb{R} \rightarrow \mathbb{R}$)

Diese “Kontraktionsbedingung” sichert $\lim_{n \rightarrow \infty} \Delta x^n = 0$, d.h. der Abstand benachbarter Iterierter konvergiert gegen null. Bedingung (*) meist nur lokal erfüllt, d.h. in einer Umgebung der Lösung $\bar{x}$ mit $f(\bar{x}) = 0$

($\rightarrow$ “lokale” Konvergenz des Newton-Verfahrens)

Beispiel Newton-Verfahren (für $f : \mathbb{R} \rightarrow \mathbb{R}$)

($f : \mathbb{R}^2 \rightarrow \mathbb{R}^2$)

$$f(x) = x^2 - 2, \text{ Startwert } x_0 = 1.5, \quad f'(x) = 2x$$

$$x_1 = x_0 - \frac{f(x_0)}{f'(x_0)} = 1.5 - \frac{0.25}{3} = \frac{3}{2} - \frac{1}{12} = \frac{17}{12}$$

$$\frac{17}{12} = 1.\overline{41}6$$

$$\sqrt{2} = 1.\overline{41}42\dots$$

Uns interessiert nun aber auch, was passiert, wenn wir in der Nähe eines Fixpunktes starten. Divergieren die Trajektorien oder bleiben sie in der Nähe, oder konvergieren sie gar gegen den Fixpunkt?

Definition 8.4.5 (Stabilität, asymptotische Stabilität eines Fixpunkts).

Ein Fixpunkt x^* eines autonomen dynamischen Systems $\dot{x} = f(x)$ heisst **stabil**, falls es für jedes $\epsilon > 0$ ein $\delta > 0$ gibt, so dass alle Trajektorien, die in der δ -Umgebung von x^* starten, in der ϵ -Umgebung bleiben, d.h. für jedes x_0 mit $\|x_0 - x^*\| \leq \delta$ gilt, dass jeder Punkt $x(t)$ der Lösungskurve $x(\cdot)$ des AWP

$$\dot{x}(t) = f(x(t)), \quad x(0) = x_0$$

die Gleichung $\|x(t) - x^*\| \leq \epsilon$ erfüllt. Ein Fixpunkt heisst **asymptotisch stabil**, wenn er stabil ist und zusätzlich gilt, dass es ein $\delta > 0$ gibt, so dass für die Lösung $x(\cdot)$ des AWP für jedes x_0 mit $\|x_0 - x^*\| \leq \delta$ gilt:

$$\lim_{t \rightarrow \infty} x(t) = x^*.$$

Wie bekommen wir heraus, ob ein Fixpunkt x^* (asymptotisch) stabil ist? In der Nähe des Fixpunktes können wir das System linearisieren, und erhalten:

$$f(x^* + \Delta x) \approx \underbrace{f(x^*)}_{=0} + \underbrace{\frac{\partial f}{\partial x}(x^*)}_{\text{Jacobi-Matrix} =: A} \cdot \Delta x = A \cdot \Delta x$$

Für $x(t) = x^* + \Delta x(t)$ gilt nun also

$$\dot{x}(t) = f(x(t)) = f(x^* + \Delta x(t)) \approx A \cdot \Delta x(t)$$

und umgekehrt gilt natürlich auch

$$\dot{x}(t) = \frac{d}{dt}(x^* + \Delta x(t)) = \frac{d(\Delta x(t))}{dt} = \Delta \dot{x}(t),$$

da x^* konstant ist. Die Abweichung $\Delta x(t)$ vom Fixpunkt gehorcht also näherungsweise der linearen autonomen Differentialgleichung

$$\Delta \dot{x}(t) = A \cdot \Delta x(t).$$

Wenn man sich die Eigenwerte von $A := \frac{\partial f}{\partial x}(x^*)$ ansieht, erfährt man oft schon viel über das System, z.B. ob es in der Nähe von x^* stabil ist, oder ob es oszilliert. Es gilt insbesondere der folgende Satz (ohne Beweis).

Satz 8.4.6 (Eigenwertkriterium für asymptotische Stabilität eines Fixpunkts).

Sei $U \subset \mathbb{R}^n$, $f \in C^1(U, \mathbb{R}^n)$, und $x^* \in U$ erfülle die Fixpunktgleichung $f(x^*) = 0$.

Wenn alle Eigenwerte der Jacobi-Matrix $\frac{\partial f}{\partial x}(x^*)$ negativen Realteil haben, ist der Fixpunkt x^* asymptotisch stabil.

Wir illustrieren den Satz an dem Insulin-Zucker-Modell aus dem Beispiel 8.0.25 zu Beginn dieses Kapitels.

Beispiel 8.4.7 (Insulin-Zucker-Modell).

Wir betrachten das Insulin-Blutzucker-Modell aus Beispiel 8.0.25 mit der Systemfunktion

$f(x) = \begin{pmatrix} x_2^2 - x_1 x_2 \\ -x_1 x_2 + 1 \end{pmatrix}$. Aus der Fixpunktgleichung $f(x^*) = 0$ finden wir

$$x^* = \begin{pmatrix} 1 \\ 1 \end{pmatrix}. \quad \left[\begin{pmatrix} -1 \\ -1 \end{pmatrix} \text{ ist zweiter, aber unphysikalischer Fixpunkt} \right]$$

Wir berechnen $\frac{\partial f}{\partial x}(x) = \begin{pmatrix} -x_2 & 2x_2 - x_1 \\ -x_2 & -x_1 \end{pmatrix},$

$$A := \frac{\partial f}{\partial x}(x^*) = \begin{pmatrix} -1 & 1 \\ -1 & -1 \end{pmatrix},$$

$$\begin{aligned} \det(A - \lambda \mathbb{I}) &= (-1 - \lambda)(-1 - \lambda) - (-1)1 \\ &= (1 + \lambda)^2 + 1 = \lambda^2 + 2\lambda + 2 \end{aligned}$$

$$\Leftrightarrow \lambda_{1,2} = -1 \pm \sqrt{1 - 2} = -1 \pm i.$$

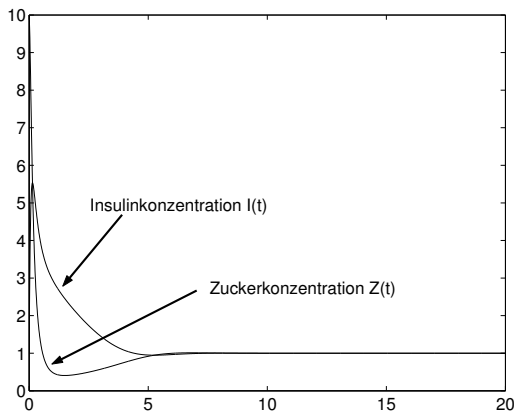


Abbildung 8.13: Langzeitverhalten des Insulin-Zucker-Systems.

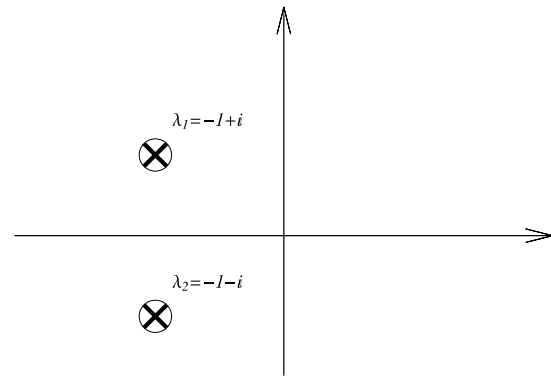
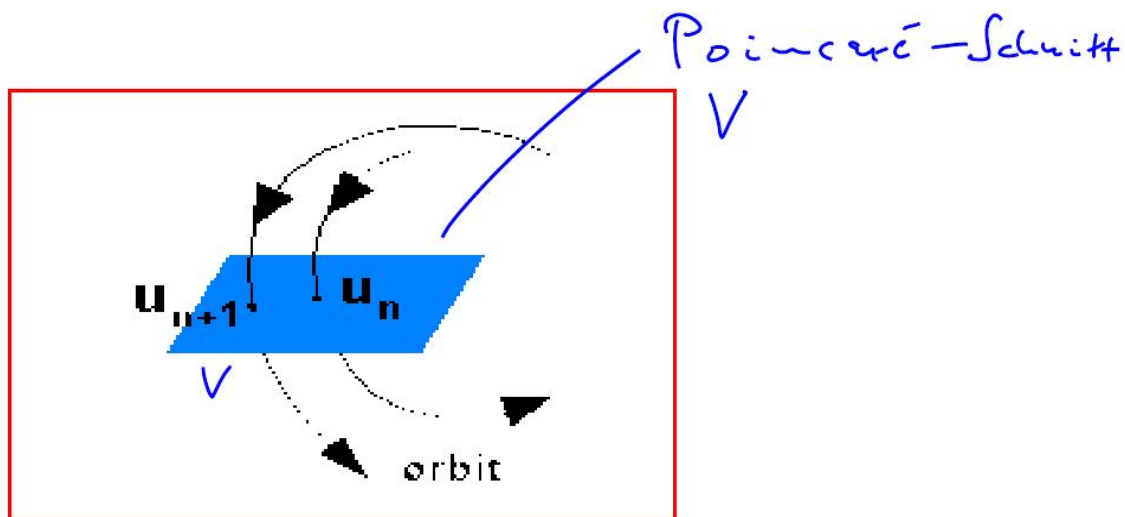


Abbildung 8.14: Eigenwerte der Jacobi-Matrix am Fixpunkt.

Aus der Tatsache, dass die Realteile beider Eigenwerte gleich -1 und damit kleiner als Null sind (in Abbildung 8.14 sind sie in die komplexe Ebene eingetragen), schliessen wir mit Satz 8.4.6, dass der Fixpunkt $x_1^* = x_2^* = 1$ asymptotisch stabil ist; zusätzlich sehen wir an den Eigenwerten noch, dass das dynamische System in der Nähe des Fixpunkts gedämpft oszilliert. Diese Oszillation ist es, die ein Unterschwingen des Blutzuckerspiegels nach einer vorherigen Erhöhung verursacht. In Abbildung 8.13 sieht man das Langzeitverhalten nach der gleichen Auslenkung wie in Abbildung 8.1.

Poincare-Abbildung



Definiere eine “Mannigfaltigkeit” (z.B. Ebene im $\mathbb{R}^3$) im Phasenraum, die transversal (lokal orthogonal) zu einem periodischem Orbit liegt (**Poincaré-Schnitt V**) und eine Abbildung

$$P : V \rightarrow V$$

$$u_n \mapsto u_{n+1}$$

mit u_{n+1} : erster Schnittpunkt der Lösung von $\dot{x} = f(x)$,
die im Punkt $u_n \in V$ startet.

Man nennt $P : V \rightarrow V$ auch eine phasenflu-vermittelte Abbildung, da der Punkt $u_n \in V$ durch den Phasenflu im “Strömungsfeld” (Richtungsfeld) der Differentialgleichung auf den Punkt $u_{n+1} \in V$ “transportiert” wird.

Definition 8.4.8. Sei $\Phi(t)$ periodische Lösung des Systems $\dot{x} = f(x)$. Man wähle eine transversale Mannigfaltigkeit V zu $\Phi(t)$ im Punkte a . Sei $P : V \rightarrow V$ die zugehörige Poincare-Abbildung. Dann heißt $\Phi(t)$ stabil, falls es für alle $\epsilon > 0$ ein $\delta(\epsilon) > 0$ gibt, so dass gilt:

$$\|x_0 - a\| \leq \delta, x_0 \in V \Rightarrow \left\| \underbrace{P \circ P \circ P \circ \dots \circ P}_{n\text{-mal}}(x_0) - a \right\| \leq \epsilon, n = 1, 2, 3, \dots$$

Definition 8.4.9. Voraussetzungen wie bei der vorherigen Definition. Dann heißt $\Phi(t)$ asymptotisch stabil, falls $\Phi(t)$ stabil ist und es ein $\delta > 0$ gibt mit:

$$\|x_0 - a\| \leq \delta, x_0 \in V \Rightarrow \lim_{n \rightarrow \infty} P^n(x_0) = a$$

Stabilitätsaussagen über periodische Orbits analog zu Fixpunkten der Poincare-Abbildung!

Stabilitätsanalyse periodischer Orbits

Sei $\phi(t)$ periodische Lösung von $\dot{x} = f(x)$, d.h. $\dot{\phi}(t) = f(\phi)$, $\phi(t+T) = \phi(t) \forall t > 0$. Setze

$$x(t) := \phi(t) + \boxed{y(t)} \quad \curvearrowright \text{Abweichungen vom periodischen Orbit}$$

Taylor-Entwicklungen 1.Ordnung von $f(x)$ “um” den periodischen Orbit $\phi(t)$:

$$\dot{x}(t) = f(x) = f(\phi + y) = f(\phi) + \frac{\partial f}{\partial x}(\phi)y + \dots$$

Einsetzen von $x(t) = \phi(t) + y(t)$: $\dot{\phi}(t) + \dot{y}(t) = f(\phi) + \frac{\partial f}{\partial x}(\phi)y + \dots$ in einer Umgebung von $\phi(t)$, also

$$\dot{y}(t) = \frac{\partial f}{\partial x}(\phi)y \quad (\text{Lineares System})$$

Stabilitätsanalyse über Eigenwerte von $\frac{\partial f}{\partial x}(\phi(t))$ möglich!

Alternative:

Betrachte Linearisierung der Poincare-Abbildung um den Punkt a (Durchstoppunkt des periodischen Orbits durch V)

$$\boxed{\frac{\partial P}{\partial x_0}(a)}$$

“Änderung” der Poincare-Abb. mit der Änderung des Anfangswertes x_0 .

Numerische Berechnung periodischer Orbits

Problemformulierung als Randwertaufgabe:

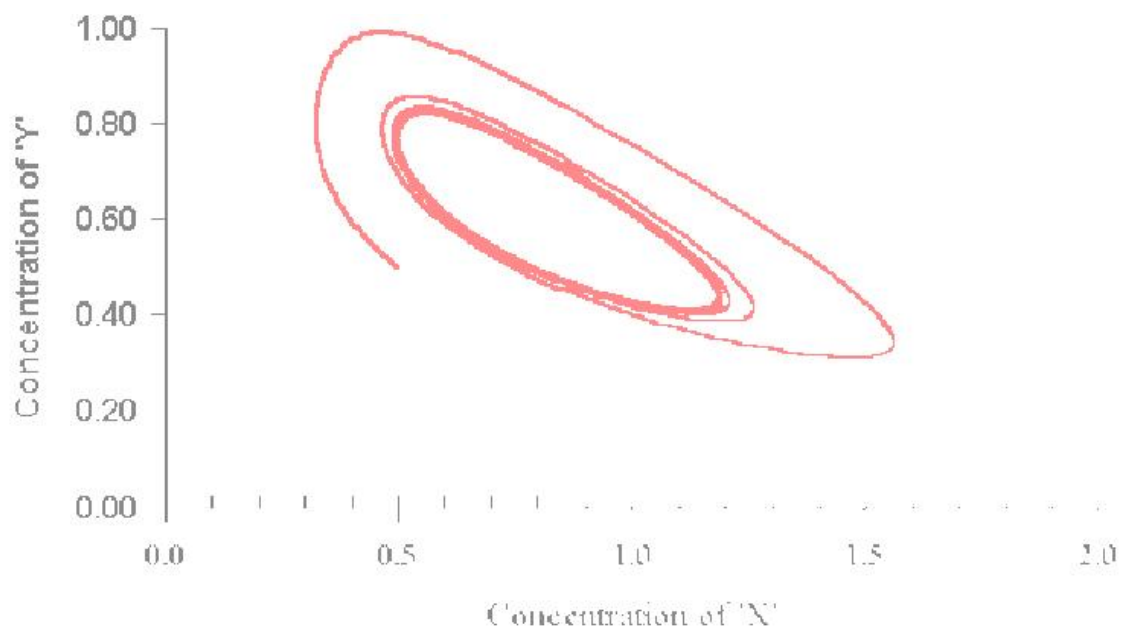
Suche Punkt $a \in \mathbb{R}^n$ mit $P(a) = a$ (P : Poincare-Abbildung)

Definiere $F(x) := P(x) - x = 0$ und suche Nullstelle von F mit Hilfe des Newton-Verfahrens !

Definition 8.4.10. Eine isolierte geschlossene Trajektorie heißt

Grenzzzyklus .

Dabei meint isoliert, dass benachbarte Trajektorien nicht geschlossen sind.



8.4.2 Bifurkationen

Heuristische Definition: Veränderungen im Auftreten und der Stabilität von Fixpunkten und periodischen Orbits.

Definition 8.4.11. Das Entstehen und Verschwinden von Fixpunkten oder die Änderung der Eigenschaften von Fixpunkten (i.e. stabil zu instabil oder umgekehrt) heisst

Soll aber für unsere

Bifurkation .

Diese Definition ist weder präzise, noch erfasst sie alle Bifurkationen.

Zwecke ausreichen.

In der Praxis:

Bifurkationen treten bei Variation von Parametern in den Systemgleichungen dynamischer Systeme auf.

Beispiel 8.4.12. Betrachte die nichtlineare Differentialgleichung

$$\dot{x} = \sin x.$$

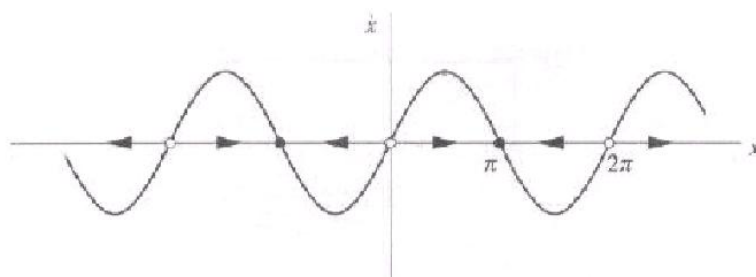
Die exakte Lösung wäre

$$t = \log \left| \frac{\csc x_0 + \cot x_0}{\csc x + \cot x} \right|, x(0) = x_0.$$

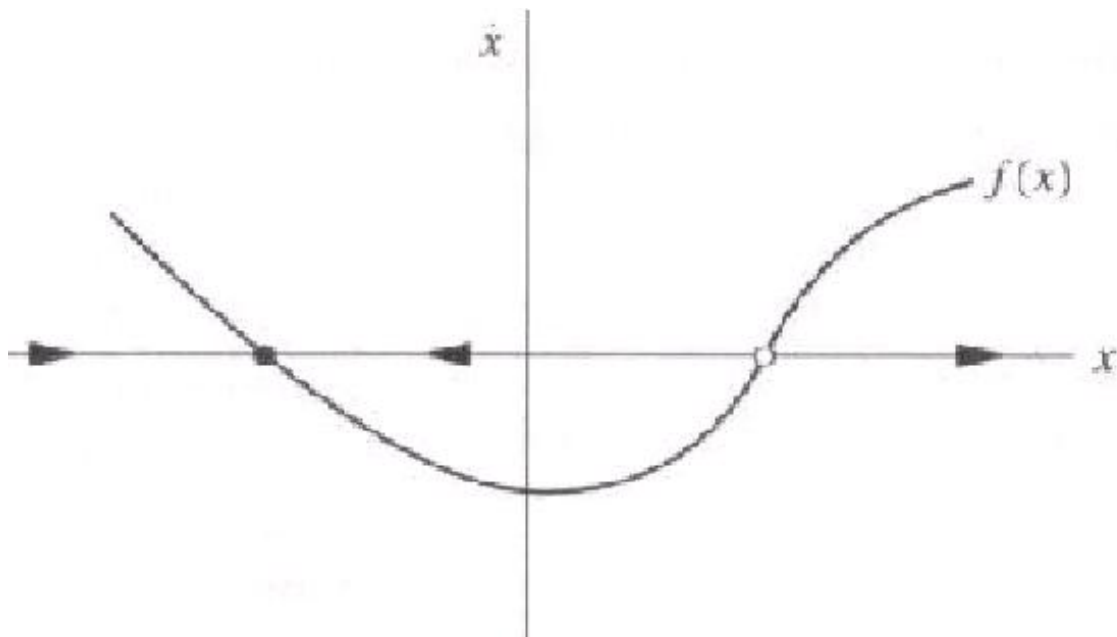
Qualitative Darstellung des Phasenflusses:

Beispiel 8.4.13.

$$\dot{x} = \sin x.$$



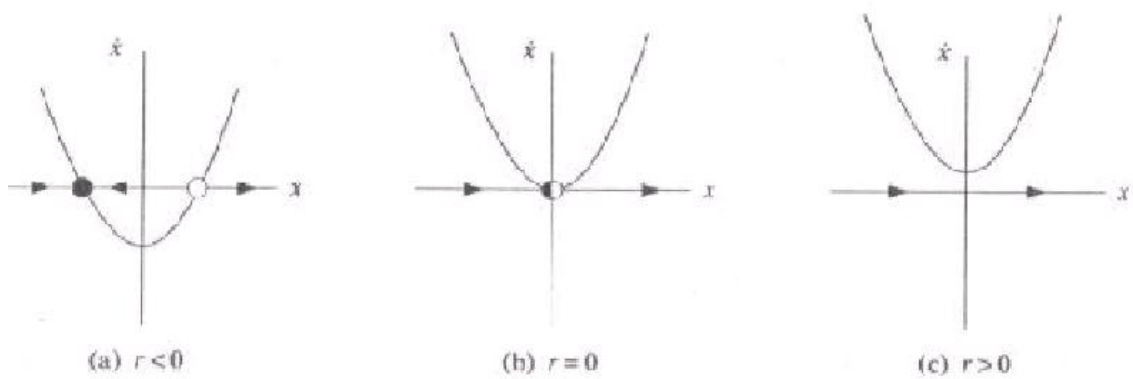
Im allgemeinen Fall $\dot{x} = f(x)$ des vorherigen Beispiels haben wir dieses Bild



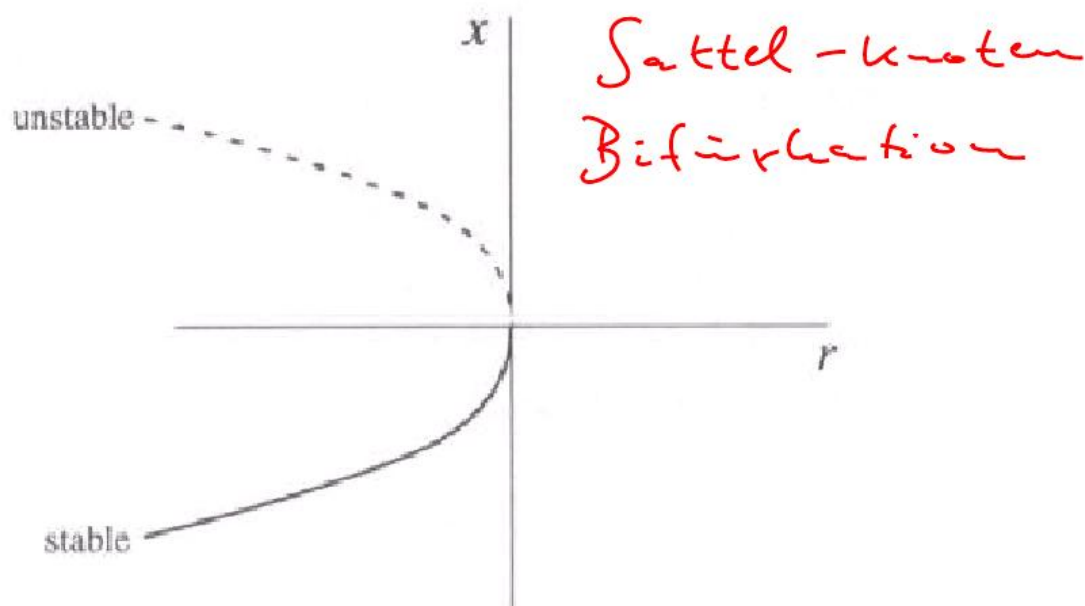
Betrachte das System erster Ordnung

$$\dot{x} = r + x^2, \quad r \in \mathbb{R}.$$

Bei Variation von r können drei Fälle eintreten.



Bifurkationsdiagramm



Darstellung der Fixpunkte x als Funktion des sog. Bifurkationsparameters r .
Betrachte wieder ein System erster Ordnung, nämlich

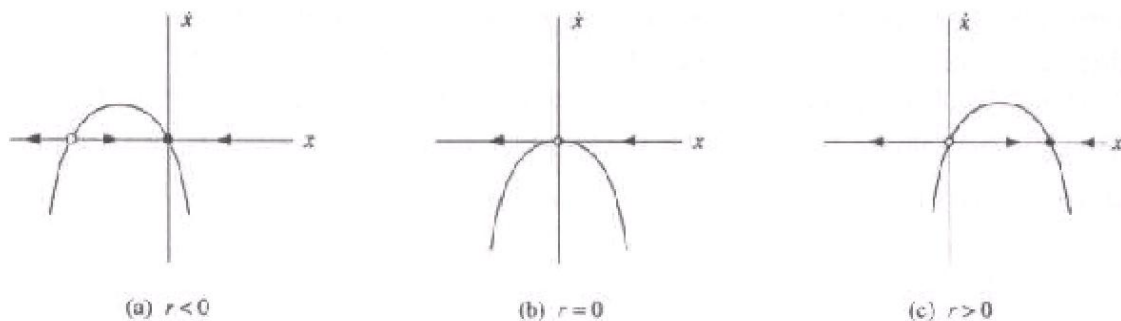
$$\dot{x} = rx - x^2, \quad r \in \mathbb{R}.$$

Diese erinnert an das logistische Wachstum.

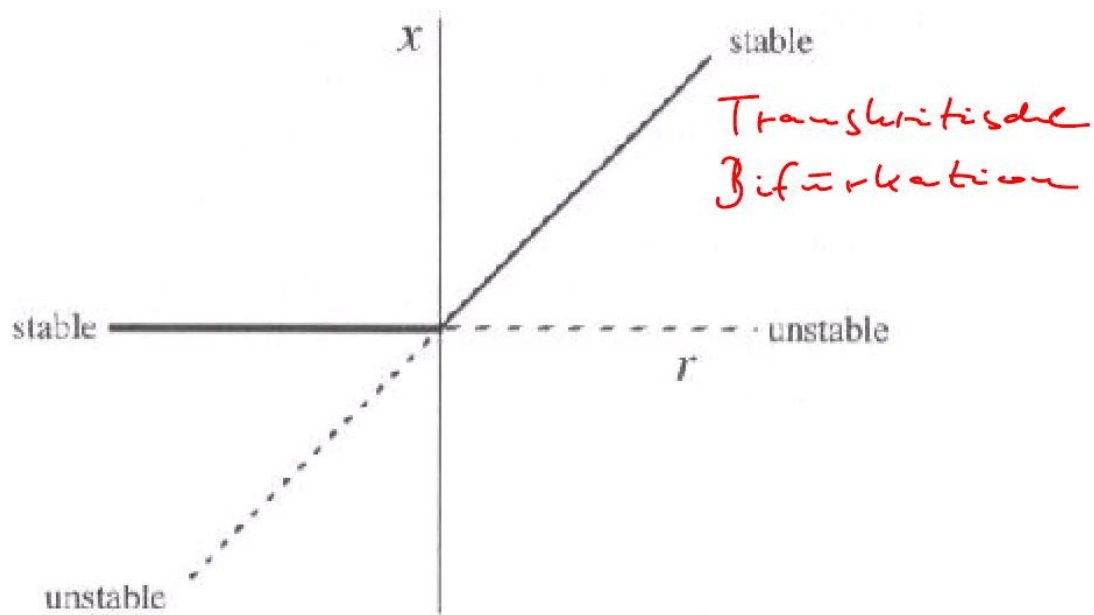
Das Wachstum von Populationen wird durch das Logistische Wachstum

$$\dot{N} = rN \left(1 - \frac{N}{K}\right)$$

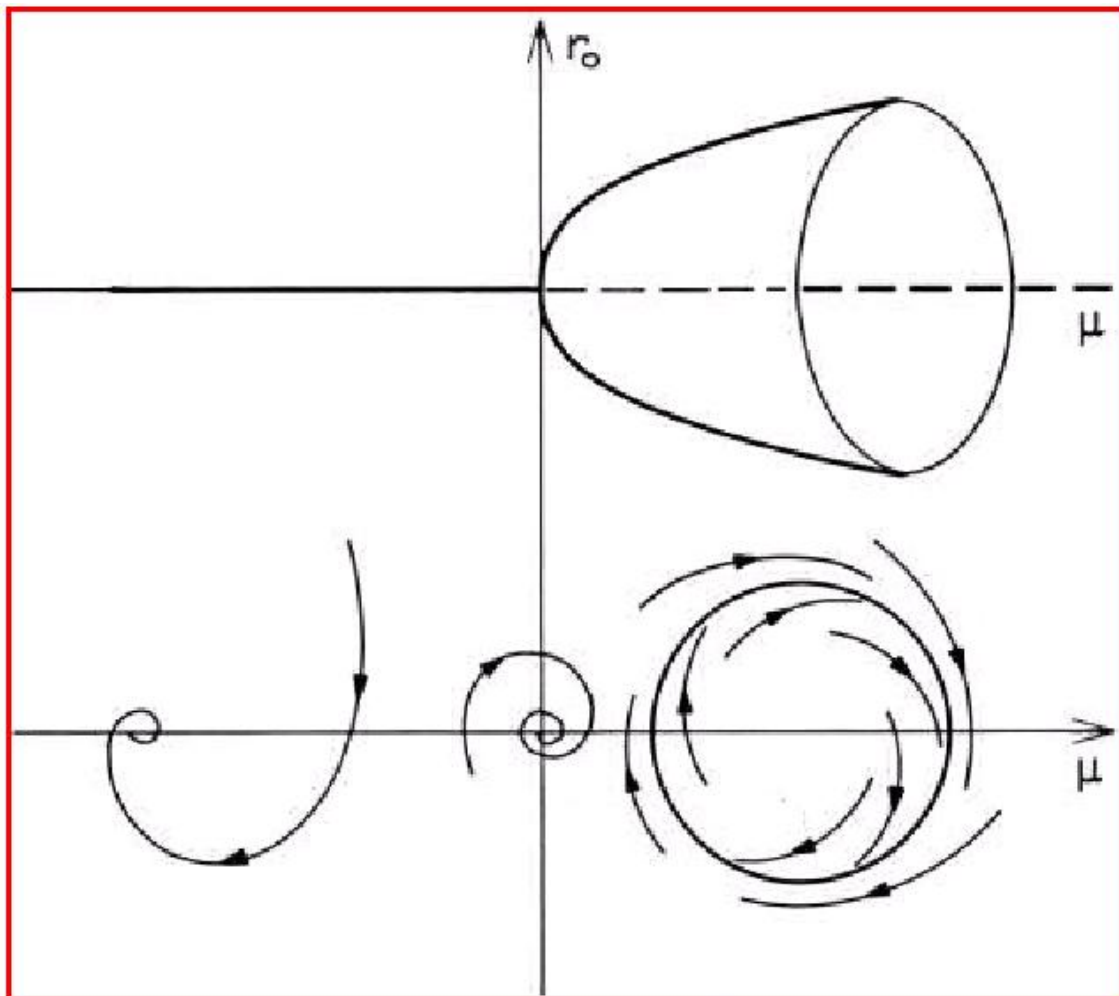
Es treten wieder drei Fälle auf:



Die Stabilität des Fixpunktes bei $x = 0$ ändert sich mit r .
Bifurkationsdiagramm



Hopf-Bifurkation: stabiler Fokus wird instabil und gleichzeitig entsteht in der Umgebung ein stabiler Grenzzyklus



Die Realteile von komplex konjugierten Eigenwerten wechseln bei der Hopf-Bifurkation das Vorzeichen.

Satz 8.4.14 (Poincare-Bendixson). Sei R eine abgeschlossene und beschränkte Teilmenge des zwei-dimensionalen Phasenraums, welche keine steady states enthält und $f(x)$ sei stetig differenzierbar. (Fixpunkte)

Gibt es eine Trajektorie $C(t)$, die ganz in R liegt (d.h. in R startet und für alle Zeiten in R bleibt), so ist $C(t)$ ein periodischer Orbit oder konvergiert für $t \rightarrow \infty$ zu einem periodischen Orbit (Grenzzyklus).

!!! Wichtiges Hilfsmittel in der “phase plane analysis” !!!

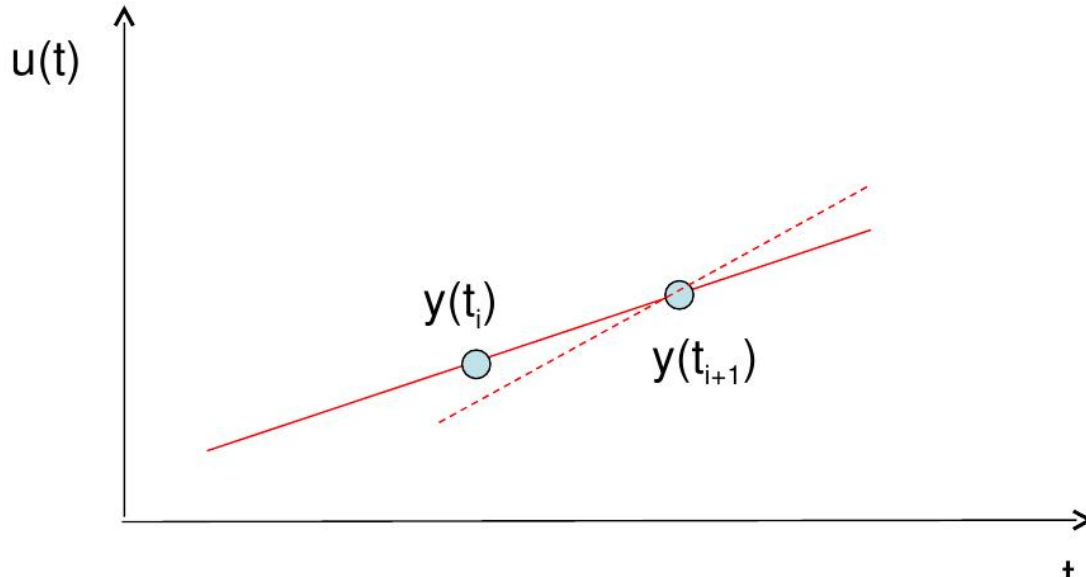
8.4.3 Numerische Lösungsverfahren

Beschränkung auf sehr einfache sog. Einschrittverfahren!

→ explizites und implizites Euler-Verfahren

Numerische Approximation von Lösungen

Hier: Beschränkung auf autonomen Fall: $du/dt = f(u)$, $df/dt = 0$



Durch die Differentialgleichung $du/dt = f(u)$ ist die Steigung (Tangente) einer Lösungskurve in jedem Punkt bekannt.

Eulersches Polygonzugverfahren: laufe auf Tangente zum nächsten Iterationspunkt eines diskreten Zeitgitters und approximiere $u(t)$ durch $y(t_i)$, $i = 0, \dots, n$, $y(t_0) = u(t_0) = u_0$ (Anfangswert)

Explizite Verfahren

Beim einfachsten *Differenzenverfahren* werden Näherungen $y_n \approx u(t_n)$ auf einem endlichen Punktgitter des Intervalls I , $t_0 < t_1 < \dots < t_n < \dots < t_N = t_0 + T$, mit Gitterweiten $h_n = t_n - t_{n-1}$ berechnet z.B. durch die rekursive Beziehung

$$h_n^{-1}\{y_n - y_{n-1}\} = f(t_{n-1}, y_{n-1}), \quad n \geq 1, \quad y_0 = u_0 \quad \text{Explizites Euler-Verfahren}$$

Begriff “explizit”: zur Berechnung eines Zeitschritts werden nur Informationen aus vorangegangenen Zeitschritten verwendet und die Berechnung des neuen Punktes kann direkt durch Additionen und Multiplikationen erfolgen:

$$y_n^h = y_{n-1}^h + h_n f(t_{n-1}, y_{n-1}^h), \quad n \geq 1 \quad (8.17)$$

Implizite Verfahren

$$h_n^{-1}\{y_n - y_{n-1}\} = f(t_n, y_n), \quad n \geq 1, \quad y_0 = u_0$$

Explizites Euler-Verfahren

Begriff “implizit”: zur Berechnung eines Zeitschritts fließen unbekannte Informationen am neu zu berechnenden Lösungspunkt ein. Daher muss i.a. ein lineares oder sogar ein nichtlineares Gleichungssystem gelöst werden, je nachdem ob die Funktion f linear oder nichtlinear ist.

Steife Systeme !!!

8.5 Zeitdiskrete dynamische Systeme

Wir wollen am Ende dieses Kapitels über dynamische Systeme noch eine eigentlich viel einfachere Art von System behandeln, nämlich *zeitdiskrete Systeme*, die nicht durch eine gewöhnliche Differentialgleichung beschrieben werden, sondern einfach nur durch eine wiederholte Anwendung einer Abbildung auf sich selbst. Wir definieren uns im folgenden für diese Systemklasse Begriffe wie Trajektorie, Fixpunkt, Stabilität.

Definition 8.5.1. Eine Iterationsvorschrift

$$x(k+1) = f(x(k)) \quad \left(\text{kurz auch } x_{\text{neu}} = f(x_{\text{alt}}) \text{ oder } x_+ = f(x) \right)$$

mit $k \in \mathbb{N}$ und einer Funktion $f : U \rightarrow U$ ($U \subset \mathbb{R}^n$) nennen wir **zeitdiskretes** System. Oft schreibt man statt $x(k)$ auch x_k .

Der Einfachheit halber betrachten wir hier nur autonome Systeme. Man könnte jedoch leicht eine Zeitabhängigkeit $x_{k+1} = f(x_k, k)$ einführen.

- Man nennt dynamische Systeme, die wie in den vorherigen Abschnitten durch gewöhnliche Differentialgleichungen beschrieben werden, auch manchmal *zeitkontinuierliche* dynamische Systeme, um sie von den zeitdiskreten Systemen zu unterscheiden.
- Achtung: Die Funktion f für zeitdiskrete Systeme ist etwas ganz anderes als die Funktion f für zeitkontinuierliche.

Beispiel 8.5.2 (Logistische Abbildung).

Die Vorschrift $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto f(x) := ax(1-x)$ nennt man die *logistische Abbildung*, sie ist durch eine Parabel beschrieben. Durch wiederholte Anwendung der Abbildung, wie in Abbildung 8.15 gezeigt, erhält man eine Folge $x_0, x_1 = f(x_0), x_2 = f(x_1), \dots$

Ursprung: Die logistische Abbildung wurde 1845 von dem belgischen Mathematiker P. F. Verhulst benutzt, um das Wachstum von Tierpopulationen mit der Größe x_k im k -ten Jahr zu beschreiben. Sein Ansatz war ein im Prinzip exponentielles Wachstumsmodell $x_{k+1} = c \cdot x_k$ mit Wachstumsrate c . Diese Wachstumsrate c wird dann aber nicht als konstant angenommen, sondern als von x_k abhängig, $c = c(x_k) = a(G - x_k)$, um zu berücksichtigen, dass die Wachstumsrate bei zu großer Population kleiner wird, mit einer Wachstumsgrenze G . Setzt man $G = 1$, erhält man die logistische Abbildung.

Bemerkung 8.5.3. Falls der **Anfangswert** $x(0) \in \mathbb{R}^n$ bekannt ist, dann auch die gesamte **Trajektorie** $x(1), x(2), \dots \in \mathbb{R}^n$

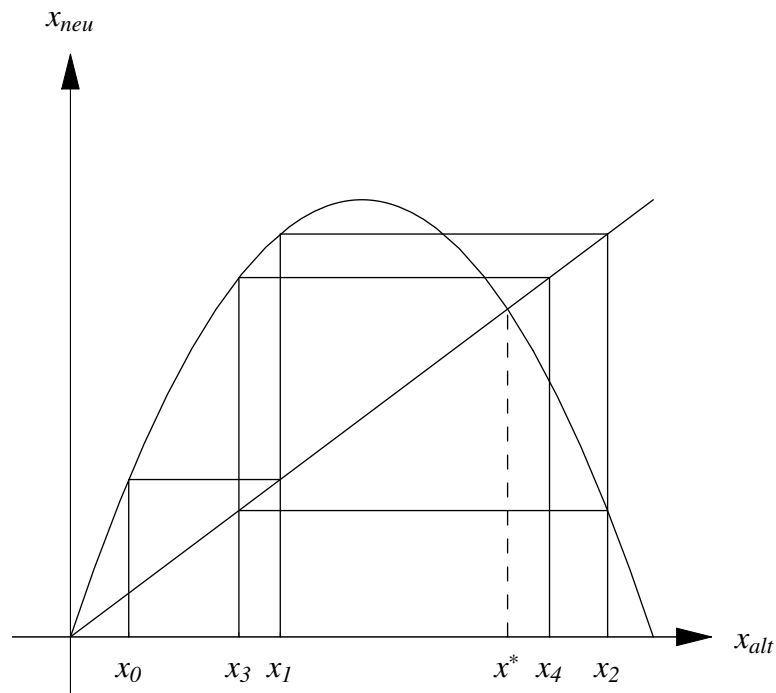


Abbildung 8.15: Einige Iterierte der logistischen Abbildung

Ein Fixpunkt ist nun etwas anders definiert als zuvor. Es soll wieder ein Punkt sein, in dem die Trajektorie verharret, wenn man in ihm startet.

Definition 8.5.4. Der Punkt $x^* \in U$ ist Fixpunkt eines zeitdiskreten autonomen Systems $x(k+1) = f(x(k))$ genau dann, wenn

$$x^* = f(x^*).$$

In Abbildung 8.15 sieht man, dass im Falle $n = 1$ ein Fixpunkt als Schnittpunkt der Winkelhalbierenden mit dem Graphen der Funktion f aufgefasst werden kann.

Beispiel 8.5.5 (Fixpunkte der logistischen Abbildung).

$$\begin{aligned} f(x) &= 2 \cdot x \cdot (1 - x) \\ x^* &= 2x^*(1 - x^*) \\ \Rightarrow x^* &= 0 \quad \text{oder} \quad x^* = \frac{1}{2} \end{aligned}$$

8.5.1 Lineare Systeme

Wichtig und besonders einfach sind auch im Zeitdiskreten die *linearen Systeme*.

Definition 8.5.6. Falls $f(x) = A \cdot x$ heisst das (autonome) zeitdiskrete System **linear**.

Beispiel 8.5.7 (Bevölkerungswachstum).

Die Bevölkerungsentwicklung eines fiktiven Landes folgt von einem Jahr zum nächsten in etwa einem linearen Gesetz der Form $x_{\text{neu}} = Ax_{\text{alt}}$, nämlich

$$\begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \\ x_6 \\ x_7 \\ x_8 \end{pmatrix}_{\text{neu}} = \begin{bmatrix} 0.88 & 0.001 & 0.06 & 0.07 & 0.01 & 0.002 & 0.001 & 0 \\ 0.098 & 0.89 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0.099 & 0.9 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0.1 & 0.9 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0.1 & 0.89 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0.099 & 0.89 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0.099 & 0.80 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0.09 & 0.75 \end{bmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \\ x_6 \\ x_7 \\ x_8 \end{pmatrix}_{\text{alt}}$$

wobei x_i die Anzahl aller Personen im i -ten Lebensjahrzehnt ist. In x_8 sind zusätzlich zu den Personen im Alter von 70-80 auch noch alle Personen über 80 Jahre enthalten.

Bedeutung: Das Feld $a_{13} = 0.06$ besagt z.B., dass jeder aus der Gruppe x_3 der 20-30 jährigen im Durchschnitt pro Jahr 0.06 Kinder bekommt, die zur Gruppe x_1 der 0-10 jährigen hinzukommen. Der nichtverschwindende Geburtenbeitrag $a_{17} = 0.001$ der 60-70 jährigen wäre auf den Beitrag von Männern mit wesentlich jüngeren Frauen zurückzuführen.

Die Felder $a_{33} = 0.9$ und $a_{43} = 0.1$ bedeuten, dass jedes Jahr 10% der Gruppe x_3 durch Älterwerden in die Gruppe x_4 übergehen. Bei den jüngeren und den älteren Jahrgängen gibt es einige Todesfälle, so dass die beiden entsprechenden Matrixeinträge sich nicht mehr zu eins summieren. Beispielsweise summieren sich $a_{11} = 0.88$ und $a_{21} = 0.098$ zu $a_{11} + a_{21} = 0.989 < 1$, und der Rest der 0-10 jährigen, d.h. 1.1%, verstirbt jedes Jahr.

Fragestellungen: Man kann sich jetzt z.B. fragen, wie sich die Altersstruktur weiterentwickelt, wenn sie in einem Jahr durch die Zahlen (in Millionen)

$$x = [2.39, 1.39, 1.02, 2.72, 4.64, 3.77, 1.73, 0.62]^T$$

gegeben ist (siehe Abbildung 8.16 ganz links).

Wie sieht z.B. die Bevölkerung (wahrscheinlich) in 20 Jahren aus, wie in 100 Jahren? Wie sah sie vermutlich vor 1 und vor 5 Jahren aus? Man erhält hierfür die einfach zu berechnenden Ausdrücke $A^{20}x$, $A^{100}x$, $A^{-1}x$ und $A^{-5}x$. In Abbildung 8.16 zeigen wir die Vektoren $x(0) = x$, $x(20) = A^{20}x$, $x(100) = A^{100}x$ und $x(500) = A^{500}x$.

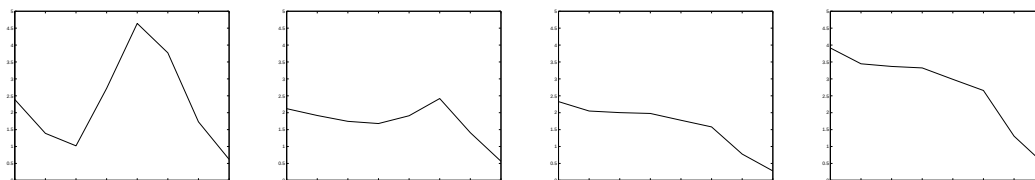


Abbildung 8.16: Bevölkerungspyramide nach 0, 20, 100 und 500 Jahren.

Tipp für Computer-Interessierte: Sie können sich die Matrix A und den Vektor x als SCILAB-Skript unter <http://www.iwr.uni-heidelberg.de/~agbock/teaching/2002ws/BIO/blatt08aufgabe4.sci> vom Netz herunterladen, dann müssen Sie die Zahlen nicht abtippen. Sodann können Sie die SCILAB-Kurzform A^n für die n -fache Matrixmultiplikation $A \cdot A \cdot \dots \cdot A$ verwenden, und den Befehl `inv(A)`, um die Inverse zu berechnen. (Was berechnet A^{-n} ?) Die entsprechende Bevölkerungsstruktur guckt man sich dann einfach mit `plot(x)` bzw. mit `plot(A^n * x)` an.

Anregung für an Modellierung Interessierte: Das Modell könnte wesentlich verbessert werden, wenn man die Bevölkerung statt in Lebensjahrzehnte in kleinere Gruppen, am besten in Lebensjahre unterteilen würde. Wie würden sich dann die Übergangszahlen von einer Gruppe zur nächsten durch Älterwerden verändern?

Oder wie könnte das Modell durch Unterscheidung in Geschlechter weiter verfeinert werden?

Stabilität und Eigenwerte

Es gilt im Zeitdiskreten ganz analog zu Satz 8.3.8:

Satz 8.5.8 (Lösung des AWP für lineare diskrete Systeme). Sei $x_+ = A \cdot x$ ein lineares System und $A \in \mathbb{R}^{n \times n}$ diagonalisierbar als

$$A = B \cdot D \cdot B^{-1}, \quad \text{mit} \quad D = \begin{pmatrix} \lambda_1 & & \\ & \ddots & \\ & & \lambda_n \end{pmatrix} \quad \text{und} \quad B = \left(v_1 \mid v_2 \mid \dots \mid v_n \right),$$

Dann ist die Trajektorie $x(0), x(1), x(2), \dots$ zu einem Anfangswert $x(0)$ durch

$$x(k) = \sum_{i=1}^n \lambda_i^k \cdot v_i \cdot a_i \quad \text{mit} \quad a = \begin{pmatrix} a_1 \\ \vdots \\ a_n \end{pmatrix} = B^{-1} \cdot x(0)$$

gegeben.

Beweis:

$$\begin{aligned} x(k) &= A^k \cdot x(0) = \underbrace{(BDB^{-1}) \cdots (BDB^{-1})}_{k\text{-mal}} x(0) \\ &= BD^k B^{-1} x(0) \\ &= B \cdot D^k \cdot a \quad \text{und} \quad D^k = D \cdot D \cdots D = \begin{pmatrix} \lambda_1^k & & 0 \\ & \ddots & \\ 0 & & \lambda_n^k \end{pmatrix} \quad \square \end{aligned}$$

Bemerkung 8.5.9. Falls $|\lambda_i| < 1$, fällt die entsprechende Komponente $\lambda_i^k \cdot v_i \cdot a_i$ ab. Falls $|\lambda_i| > 1$ wächst sie.

Achtung: Während bei zeitkontinuierlichen linearen Systemen die Stabilitätsgrenze durch die imaginäre Achse gegeben war, ist es im Zeitdiskreten der Einheitskreis in der komplexen Ebene.

Bemerkung 8.5.10. Nach langer Zeit dominiert die Komponente zum betragsgrößten Eigenwert, ähnlich wie für zeitkontinuierliche Systeme in 8.3.1, Gleichung (8.16).

Beispiel 8.5.11 (Eigenwertanalyse des Bevölkerungsmodells).

Eine Eigenwertanalyse der Systemmatrix A aus dem Bevölkerungsmodell ergibt, dass der größte Eigenwert hier $\lambda_1 = 1.0013$ ist. Da dies leicht positiv ist, wächst die Bevölkerung exponentiell, wenn auch sehr langsam. Nach langer Zeit setzt sich dabei die zugehörige Komponente durch, die proportional zum Eigenvektor v_1 ist, den wir in Abbildung 8.17 zeigen. Ein Vergleich dieses Vektors mit der Bevölkerung nach 100 oder nach 500 Jahren in Abbildung 8.16 zeigt, dass sich die zugehörige Komponente tatsächlich durchgesetzt hat.

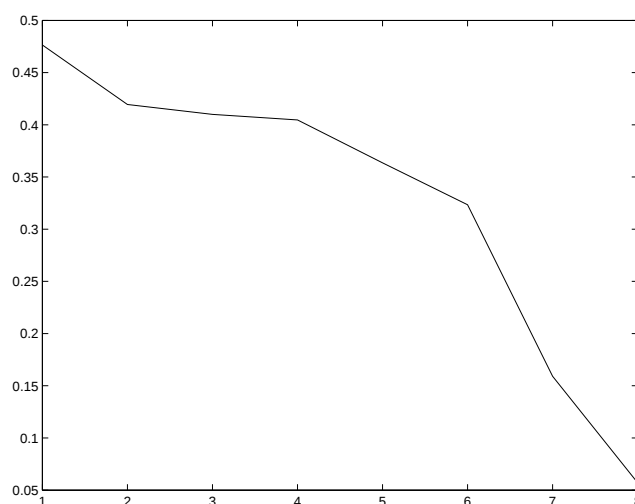


Abbildung 8.17: Der Eigenvektor v_1 zum betragsgrößten Eigenwert λ_1 des Bevölkerungsmodells. Vgl. Abbildung 8.16

8.5.2 Nichtlineare Systeme

Allgemein sind nichtlineare zeitdiskrete Systeme ebenso schwierig zu analysieren wie nichtlineare gewöhnliche Differentialgleichungen. Auch zeitdiskrete Systeme können *deterministisches Chaos* produzieren. (Um sich davon zu überzeugen, iteriere man einfach die logistische Abbildung $x_{k+1} = ax_k(1 - x_k)$ mit $a = 3.57$ für einige Zeit mit dem Anfangswert $x_0 = 0.5$, und vergleiche die gewonnene Trajektorie mit der zum Anfangswert $x_0 = 0.5001$.) Glücklicherweise kann man aber immerhin durch Linearisierung des Systems an einem Fixpunkt herausbekommen, ob der Fixpunkt stabil ist oder nicht. Dies ist die einzige Technik zur Analyse nichtlinearer zeitdiskreter Systeme, die wir hier besprechen wollen.

Definition 8.5.12 (Stabilität eines Fixpunkts).

Ein Fixpunkt x^* eines zeitdiskreten dynamischen Systems $x(k+1) = f(x(k))$ ist **asymptotisch stabil**, wenn alle Folgen $x(0), x(1), \dots$, die mit $x(0)$ in einer Umgebung von x^* starten, gegen x^* konvergieren:

$$\exists \epsilon > 0 \quad \forall x(0) : \quad \|x(0) - x^*\| \leq \epsilon \Rightarrow \lim_{k \rightarrow \infty} x(k) = x^*.$$

Satz 8.5.13 (Stabilitätskriterium eines Fixpunkts).

Sei $f \in C^1(U, U)$, $U \subset \mathbb{R}^n$ und $x^* = f(x^*)$, und $A = \frac{\partial f}{\partial x}(x^*)$ die Jacobi-Matrix von f am Fixpunkt. Dann ist x^* asymptotisch stabil, falls der betragsgrößte Eigenwert λ_1 von A einen Betrag $|\lambda_1| < 1$ hat. Falls $|\lambda_1| > 1$, so ist x^* instabil.

Beispiel 8.5.14 (Eigenwertanalyse der Fixpunkte der logistischen Abbildung).

Wir illustrieren den Satz an den zwei Fixpunkten der logistischen Gleichung, mit $a = 2$. Da

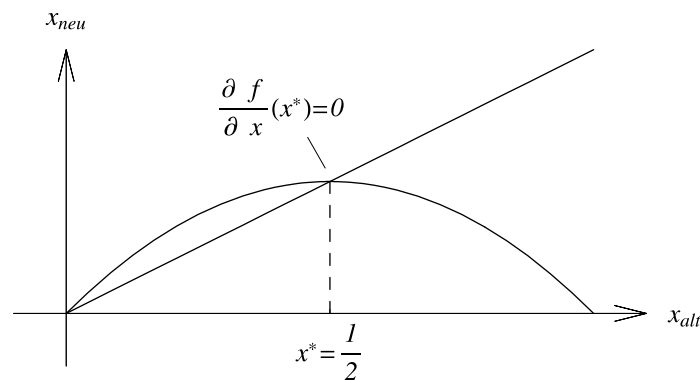


Abbildung 8.18: Stabilität des Fixpunkts $x^* = \frac{1}{2}$ der logistischen Abbildung mit $a = 2$.

die Jacobi-Matrix hier eine Zahl ist, ist sie trivialerweise gleich ihrem betragsgrößten Eigenwert, d.h. $\lambda_1 = \frac{\partial f}{\partial x}(x^*)$. Wir untersuchen die beiden Fixpunkte $x^* = 0$ und $x^* = \frac{1}{2}$ der logistischen Abbildung aus Beispiel 8.5.5. Für eine Illustration siehe Abbildung 8.18.

$$\begin{aligned} f(x) &= 2 \cdot x(1-x) & \frac{\partial f}{\partial x}(x) &= 2(1-x) - 2x = 2 - 4x \\ \frac{\partial f}{\partial x}(0) &= 2. & \text{Der Fixpunkt } x^* = 0 &\text{ ist instabil, da } |2| > 1. \\ \frac{\partial f}{\partial x}\left(\frac{1}{2}\right) &= 2 - 4 \cdot \frac{1}{2} = 0. & \text{Der Fixpunkt } x^* = \frac{1}{2} &\text{ ist stabil, da } |0| < 1. \end{aligned}$$

Kapitel 9

Wahrscheinlichkeitstheorie

Die Wahrscheinlichkeitstheorie ist nicht nur ein Hilfsmittel für erfolgreiche Glücksspieler, sondern auch die unentbehrliche Grundlage für das Verständnis der Statistik, die für sie noch eine große Bedeutung bekommen wird. Deshalb widmen wir ihr in unserem Kurs ein ganzes Kapitel. Als Vorlage für den Aufbau dieses Kapitels diente [Kre02], aus dem wir viele Definitionen, Sätze etc. übernommen haben. Eine elementare Einführung in die Wahrscheinlichkeitsrechnung bietet z.B. [Bos99].

9.1 Endliche Wahrscheinlichkeitsräume

Wir betrachten folgendes Experiment: Eine Münze wird geworfen. Das Ergebnis sei entweder „Kopf“ oder „Zahl“. Der Ausgang eines solchen Experimentes ist nicht exakt vorraussagbar. Man müßte ein exaktes physikalisches *Modell* und alle nötigen Parameter, Anfangs- und Randdaten haben, was aber unmöglich ist. Im betrachteten Fall sprechen wir von einem **Zufallsexperiment**. Die **Wahrscheinlichkeitstheorie** analysiert Gesetzmäßigkeiten solcher Zufallsexperimente. Jeder hat eine gewisse Vorstellung von der Aussage: „Bei einer fairen Münze ist die **Wahrscheinlichkeit** für ‚Kopf‘ genauso groß wie für ‚Zahl‘.“ Intuitiv denkt man dabei etwa: „Wenn man die Münze oft (hintereinander) wirft, so konvergiert die **relative Häufigkeit** von ‚Kopf‘ (von ‚Zahl‘) gegen 1/2.“ Eine Definition der Wahrscheinlichkeit mit Hilfe der relativen Häufigkeiten ist im Allgemeinen jedoch problematisch. Mathematiker definieren daher lieber abstrakt einen Wahrscheinlichkeitsbegriff und stellen dann anschließend einen Zusammenhang zwischen Wahrscheinlichkeitswert und relativer Häufigkeit her (s. Satz 9.1.55). In einigen anwendungsorientierten Beispielen werden wir uns aber zum besseren Verständnis Wahrscheinlichkeiten durch relative Häufigkeiten definieren.

Beispiel 9.1.1 (Zweimaliges Würfeln).

Experiment: Es wird zweimal hintereinander gewürfelt. Die Menge aller möglichen Kombinationen ist

$$\Omega := \{(i, j) | 1 \leq i, j \leq 6\}.$$

Also gibt es $|\Omega| = 36$ mögliche Ausgänge des Experimentes. Bei einem sogenannten **fairen Würfel** sind alle diese Ausgänge (**Elementarereignisse**) gleichwahrscheinlich. Z.B. geschieht

das Ereignis $\{(1, 2)\}$ = „erst 1, dann 2“ mit einer Wahrscheinlichkeit von $1/36$. Das Ereignis „Summe der Augenzahlen ist höchstens 3“ entspricht der Menge $A := \{(1, 1), (1, 2), (2, 1)\}$. Es gilt also $|A| = 3$ und somit ist die Wahrscheinlichkeit für dieses Ereignis gleich $3/36 = 1/12$.

9.1.1 Elementare Definitionen

Definition 9.1.2 (Endlicher Wahrscheinlichkeitsraum).

Sei Ω eine nicht-leere, endliche Menge, also o.b.d.A. $\Omega = \{1, 2, \dots, N\}$ und $\mathcal{P}(\Omega)$ deren Potenzmenge, d.h. die Menge aller Teilmengen von Ω .

1. Eine **Wahrscheinlichkeitsverteilung** (oder auch ein **Wahrscheinlichkeitsmaß**) auf Ω ist eine Abbildung $P : \mathcal{P}(\Omega) \rightarrow [0, 1]$ mit folgenden Eigenschaften:

$$P(\Omega) = 1, \quad (9.1)$$

$$P(A \cup B) = P(A) + P(B) \quad \text{für } A \cap B = \emptyset. \quad (9.2)$$

Die Menge Ω nennen wir **Ergebnismenge** oder auch **Ergebnisraum**.

2. Teilmengen $A \subset \Omega$ heißen **Ereignisse**, $P(A)$ heißt **Wahrscheinlichkeit von A** .
3. Eine Menge $\{\omega\}$ mit $\omega \in \Omega$ heißt **Elementarereignis**.
4. Das Paar (Ω, P) heißt **Wahrscheinlichkeitsraum** (genauer: endlicher Wahrscheinlichkeitsraum).
5. Wir nennen Ω das **sichere Ereignis** und $\emptyset$ das **unmögliche Ereignis**.

Bemerkung 9.1.3 (Wahrscheinlichkeitsmaß als Voraussage).

Auch wenn wir hier, wie angekündigt, mathematisch vorgehen und Wahrscheinlichkeiten von Ereignissen durch eine abstrakt gegebene Funktion P definieren, ohne dies weiter zu erklären, sollte jeder eine intuitive Vorstellung von Wahrscheinlichkeit haben. Das Wahrscheinlichkeitsmaß können wir auch als **Voraussage** über die möglichen Ausgänge eines Zufallsexperimentes interpretieren. Eine solche Sichtweise wird z.B. das Verständnis des Begriffes der *bedingten Wahrscheinlichkeit* (s. Kapitel 9.1.2) unterstützen.

Satz 9.1.4 (Eigenschaften eines Wahrscheinlichkeitsmaßes).

Seien (Ω, P) ein endlicher Wahrscheinlichkeitsraum und $A, B \in \mathcal{P}(\Omega)$. Es gilt:

- 1.

$$P(A^c) = 1 - P(A),$$

wobei $A^c = \Omega \setminus A$ das **Komplement** von A ist. Speziell gilt

$$P(\emptyset) = 0.$$

- 2.

$$A \subset B \Rightarrow P(A) \leq P(B).$$

3.

$$P(A \setminus B) = P(A) - P(A \cap B).$$

4. Falls $A_1, \dots, A_n$ paarweise disjunkt sind, d.h. für $i \neq j$ gilt $A_i \cap A_j = \emptyset$, dann gilt

$$P\left(\bigcup_{i=1}^n A_i\right) = \sum_{i=1}^n P(A_i).$$

Speziell gilt

$$P(A) = \sum_{\omega \in A} P(\{\omega\}).$$

5. Für beliebige (i.a. nicht paarweise disjunkte) $A_1, \dots, A_n \in \mathcal{P}(\Omega)$ gilt

$$P\left(\bigcup_{i=1}^n A_i\right) \leq \sum_{i=1}^n P(A_i).$$

6.

$$P(A \cup B) = P(A) + P(B) - P(A \cap B).$$

Definition 9.1.5 (Wahrscheinlichkeitsfunktion).

Die Abbildung

$$P : \Omega \rightarrow [0, 1], \quad (9.3)$$

$$\omega \mapsto P(\{\omega\}) =: P(\omega). \quad (9.4)$$

heißt **Wahrscheinlichkeitsfunktion**. Diese bezeichnen wir ebenfalls mit P . Aus dem jeweiligen Zusammenhang sollte hervorgehen, ob mit P das Wahrscheinlichkeitsmaß oder die Wahrscheinlichkeitsfunktion gemeint ist.

Bemerkung 9.1.6. (Zusammenhang zwischen Wahrscheinlichkeitsmaß und Wahrscheinlichkeitsfunktion)

Bei einem endlichen Wahrscheinlichkeitsraum ist auch umgekehrt das Wahrscheinlichkeitsmaß durch die Wahrscheinlichkeitsfunktion bestimmt. Dies gilt auch noch für abzählbare Wahrscheinlichkeitsräume (s. Kapitel 9.2.1).

Definition 9.1.7 (Laplacescher Wahrscheinlichkeitsraum).

Sei (Ω, P) endlicher Wahrscheinlichkeitsraum. Falls alle Elementarereignisse die gleiche Wahrscheinlichkeit haben, heißt P **Gleichverteilung**, und (Ω, P) heißt **Laplacescher Wahrscheinlichkeitsraum**. Es gilt dann:

$$P(\omega) = \frac{1}{|\Omega|} \quad \text{für alle } \omega \in \Omega, \quad (9.5)$$

$$P(A) = \frac{|A|}{|\Omega|} \quad \text{für } A \subset \Omega, \quad (9.6)$$

wobei $|\Omega|$, $|A|$ die Anzahl der Elemente in Ω bzw. A ist.

Beispiel 9.1.8 („6 Richtige im Lotto 6 aus 49“).

Wir berechnen die Wahrscheinlichkeit dafür, dass 6 bestimmte Zahlen (der eigene Tipp) zufällig als Gewinnzahlen gezogen werden, auf zwei verschiedene Weisen. Unser Tipp bestehe aus den sechs verschiedenen Zahlen $t_1, \dots, t_6$.

1. Als Ergebnismenge Ω_1 nehmen wir hier die Menge aller *sechs-elementigen Teilmengen* der Menge $\{1, \dots, 49\}$. Wir unterscheiden also nicht, in welcher Reihenfolge die Zahlen gezogen werden.

$$\begin{aligned} \Omega_1 = \{ \{w_1, \dots, w_6\} \mid w_i \in \{1, \dots, 49\} \text{ für alle } 1 \leq i \leq 6 \\ \text{und } w_i \neq w_j \text{ für } i \neq j \text{ und } 1 \leq i, j \leq 6 \} \end{aligned}$$

Die Anzahl dieser Teilmengen ist

$$|\Omega_1| = \binom{49}{6} = 13983816. \quad (9.7)$$

Jede Ziehung (jedes Elementarereignis) habe den gleichen Wahrscheinlichkeitswert, insbesondere auch das Elementarereignis $A_1 := \{t_1, \dots, t_6\}$, das unserem Tipp entspricht. Also

$$P_1(A_1) = \frac{1}{|\Omega|} \approx 7.1511 \cdot 10^{-8}.$$

2. Jetzt nehmen wir als Elementarereignisse alle *Sechsertupel* von paarweise verschiedenen ganzen Zahlen zwischen 1 und 49. Es kommt also auf die Reihenfolge bei der Ziehung an. Z.B. sind die Tupel $(1, 2, 3, 4, 5, 6)$ und $(6, 5, 4, 3, 2, 1)$ voneinander verschieden.

$$\begin{aligned} \Omega_2 = \{ (w_1, \dots, w_6) \mid w_i \in \{1, \dots, 49\}, \text{ für alle } 1 \leq i \leq 6, \\ w_i \neq w_j \text{ für } i \neq j \text{ und } 1 \leq i, j \leq 6 \}. \end{aligned}$$

Die Anzahl solcher Sechsertupel ist

$$\begin{aligned} |\Omega_2| &= 49 \cdot 48 \cdots 44 \\ &= \frac{49!}{43!}. \end{aligned}$$

Das Ereignis „6 Richtige“ entspricht der Menge

$$A_2 := \{ (\omega_1, \dots, \omega_6) \mid \{\omega_1, \dots, \omega_6\} = \{t_1, \dots, t_6\} \}.$$

Die Menge A_2 besteht also gerade aus allen Sechsertupeln, die aus $(t_1, \dots, t_6)$ durch Permutation hervorgehen. Für den Lottogewinn ist es ja egal, in welcher Reihenfolge die Ge-

winnzahlen gezogen werden. Es gilt also $|A_2| = 6!$. Wir erhalten also

$$\begin{aligned} P_2(A_2) &= \frac{|A_2|}{|\Omega_2|} \\ &= \frac{6! (49 - 6)!}{49!} \\ &= \frac{1}{\binom{49}{6}} \\ &\approx 7.1511 \cdot 10^{-8}, \end{aligned}$$

also letztlich das gleiche Ergebnis wie bei der ersten Rechnung.

Beispiel 9.1.9 (Dreimal Würfeln mit Laplace-Würfel).

Wie groß ist die Wahrscheinlichkeit dafür, dass dabei keine Wiederholung vorkommt? Wir wählen

$$\Omega = \{(w_1, w_2, w_3) \mid w_i \in \{1, 2, 3, 4, 5, 6\} \text{ für } 1 \leq i \leq 3\}$$

als Ergebnismenge. Die Anzahl aller möglichen Elementarereignisse (Dreiertupel) ist 6^3 . Das Ereignis „keine Wiederholung“ entspricht der Menge A aller Dreiertupel, in denen alle drei Zahlen verschieden sind. Es gibt genau $6 \cdot 5 \cdot 4 = \frac{6!}{3!}$ solche Dreiertupel. Also ist

$$P(A) = \frac{6 \cdot 5 \cdot 4}{6^3} = \frac{5}{9}.$$

9.1.2 Bedingte Wahrscheinlichkeit

In Bemerkung 9.1.3 hatten wir schon erwähnt, dass man ein gegebenes Wahrscheinlichkeitsmaß als Voraussage für ein Zufallsexperiment interpretieren kann. Wenn man nun zusätzliche Informationen über das Experiment erhält, so kann man diese Voraussage „verbessern“. Z.B. hat man nach einem einfachen Experiment wie Münzwurf die Information, wie das Experiment ausgefallen ist, und man kann mit dieser vollständigen Information im Nachhinein sogar eine deterministische „Voraussage“ (die dann ihren Namen eigentlich nicht mehr verdient) machen, d.h. man wird nicht mehr das a priori gegebene Wahrscheinlichkeitsmaß betrachten, sondern vielmehr ein anderes (deterministisches), das jedem Ereignis entweder die Wahrscheinlichkeit 0 oder 1 zuordnet. Im allgemeinen erhält man keine vollständige Information, sondern nur eine solche der Art, dass bestimmte Ereignisse sicher eintreten. Dementsprechend geht man zu einem neuen Wahrscheinlichkeitsmaß über.

Ein weiteres Beispiel ist die Wahrscheinlichkeit für den Erfolg bei einer bestimmten medizinischen Operation. Diese ist üblicherweise über die relative Häufigkeit „Anzahl der Erfolge geteilt durch Gesamtzahl der Operationen“ definiert. Bei zusätzlicher Information über den Patienten, z.B. über dessen Alter, erscheint es sinnvoll, dieses bei der Voraussage zu berücksichtigen und z.B. die Erfolgswahrscheinlichkeit durch die relative Häufigkeit innerhalb der Altersklasse des Patienten zu definieren.

Beispiel 9.1.10. (Voraussage für den zweifachen Münzwurf bei zusätzlicher Information)

Wir betrachten zwei aufeinanderfolgende Münzwürfe mit einer fairen Münze. Wie groß ist die Wahrscheinlichkeit dafür, dass „zweimal Kopf“ fällt (Ereignis A), wenn man weiß, dass

1. Fall: der erste Wurf das Ergebnis „Kopf“ hat (Ereignis B_1).
2. Fall: mindestens ein Wurf gleich „Kopf“ ist (Ereignis B_2).

Als Ergebnisraum wählen wir

$$\Omega := \{(K, K), (K, Z), (Z, K), (Z, Z)\}.$$

Da wir die Münze als fair annehmen, hat jedes Elementarereignis die Wahrscheinlichkeit $1/4$. Für unsere speziell betrachteten Ereignisse gilt

$$\begin{aligned} A &= \{(K, K)\}, \\ P(A) &= \frac{1}{4}, \\ B_1 &= \{(K, K), (K, Z)\}, \\ P(B_1) &= \frac{1}{2}, \\ B_2 &= \{(K, K), (K, Z), (Z, K)\}, \\ P(B_2) &= \frac{3}{4}. \end{aligned}$$

1. Fall: Aufgrund der zusätzlichen Informationen, dass das Ereignis B_1 eintritt, können die Elementarereignisse (Z, Z) und (Z, K) völlig ausgeschlossen werden. Es können also nur (K, K) oder (K, Z) eintreten. Ohne jegliche weitere Information sind diese beiden als gleichwahrscheinlich anzunehmen. Durch diese Überlegungen ordnen wir insbesondere dem Ereignis (K, K) eine neue Wahrscheinlichkeit zu:

$$P(A|B_1) = \frac{1}{2}.$$

Wir bezeichnen diese als **die bedingte Wahrscheinlichkeit des Ereignisses (K, K) bei gegebenem B_1** .

2. Fall: Es können nur $(K, K), (K, Z), (Z, K)$ eintreten. Wieder sehen wir diese Elementarereignisse als gleichwahrscheinlich an. Also

$$P(A|B_2) = \frac{1}{3}.$$

In beiden Fällen werden die möglichen Elementarereignisse auf eine Menge $B_i \subset \Omega$ reduziert. Wie wir sehen, ist die bedingte Wahrscheinlichkeit für das Ereignis A bei gegebenem B gleich

$$\begin{aligned} P(A|B) &= \frac{|A \cap B|}{|B|} \\ &= \frac{P(A \cap B)}{P(B)}. \end{aligned}$$

Mit Hilfe des letzten Ausdrucks definieren wir allgemein die bedingte Wahrscheinlichkeit.

Definition 9.1.11 (Bedingte Wahrscheinlichkeit).

Seien (Ω, P) ein endlicher Wahrscheinlichkeitsraum, $B \subset \Omega$ mit $P(B) > 0$ und $A \in \Omega$. Die **bedingte Wahrscheinlichkeit von A bei gegebenen B** ist

$$P(A|B) := \frac{P(A \cap B)}{P(B)}. \quad (9.8)$$

Bemerkung 9.1.12. Es folgt

$$P(A \cap B) = P(B) \cdot P(A|B). \quad (9.9)$$

Satz 9.1.13 (zur bedingten Wahrscheinlichkeit).

Sei (Ω, P) ein endlicher Wahrscheinlichkeitsraum.

1. (Die bedingte Wahrscheinlichkeit ist ein Wahrscheinlichkeitsmaß)

Sei $P(B) > 0$. Durch

$$P_B(A) := P(A|B) \quad (9.10)$$

ist ein Wahrscheinlichkeitsmaß auf Ω definiert. Ist $A \subset B^c$ oder $P(A) = 0$, so ist $P(A|B) = 0$.

2. (Formel der totalen Wahrscheinlichkeit)

Sei $\Omega = \bigcup_{i=1}^n B_i$ mit $B_i \cap B_j = \emptyset$ für $i \neq j$ (disjunkte Zerlegung von Ω). Dann gilt für jedes $A \subset \Omega$:

$$P(A) = \sum_{\substack{1 \leq k \leq n, \\ P(B_k) > 0}} P(B_k) \cdot P(A|B_k). \quad (9.11)$$

Daher wird über alle Indizes k summiert, für die $P(B_k) > 0$. Wir schreiben der Kürze halber auch „ $\sum_{k=1}^n$ “ anstatt „ $\sum_{\substack{1 \leq k \leq n, \\ P(B_k) > 0}}$ “, wobei wir im Fall $P(B_k) = 0$ das Produkt als 0 definieren.

3. (Formel von Bayes)

Sei neben den Voraussetzungen in (2.) zusätzlich noch $P(A) > 0$ erfüllt. Dann gilt für jedes $1 \leq i \leq n$:

$$P(B_i|A) = \frac{P(B_i) \cdot P(A|B_i)}{\sum_{k=1}^n P(B_k) \cdot P(A|B_k)}. \quad (9.12)$$

Beweis:

1. Die Funktion $P : \mathcal{P}(\Omega) \rightarrow \mathbb{R}$ nimmt wegen $P(B) > 0$ und $P(A \cap B) \geq 0$ nur nicht-negative Werte an. Es gilt $P_B(\Omega) = \frac{P(B)}{P(B)} = 1$, d.h. Axiom (9.1) ist für P_B erfüllt. Für beliebige disjunkte $A_1, A_2 \subset \Omega$ („disjunkt“ heisst, dass $A_1 \cap A_2 = \emptyset$) gilt

$$\begin{aligned} P_B(A_1 \cup A_2) &= \frac{P((A_1 \cup A_2) \cap B)}{P(B)} \\ &= \frac{P((A_1 \cap B) \cup (A_2 \cap B))}{P(B)} \\ &= \frac{P(A_1 \cap B)}{P(B)} + \frac{P(A_2 \cap B)}{P(B)} \quad \text{wegen (9.2))} \\ &= P_B(A_1) + P_B(A_2), \end{aligned}$$

und es folgt Axiom (9.2) für P_B . Aus diesem folgt insbesondere für beliebiges $A \subset \Omega$, dass

$$\begin{aligned} P_B(A) &\leq P_B(A) + P_B(A^C) \quad (\text{wegen } P_B(A^C) \geq 0) \\ &= P_B(\Omega) \\ &= 1, \end{aligned}$$

womit wir nachträglich auch gezeigt haben, dass P_B keine Werte größer als 1 annimmt. Damit sind alle zu zeigenden Eigenschaften bewiesen.

2. Falls $i \neq j$, sind die Mengen $A \cap B_i$ und $A \cap B_j$ disjunkt. Außerdem gilt $A = \bigcup_k (A \cap B_k)$. Mit Hilfe von Satz 9.1.4.4 und (9.9) erhalten wir

$$\begin{aligned} P(A) &= \sum_{k=1}^n P(A \cap B_k) \\ &= \sum_{k=1}^n P(B_k) \cdot P(A|B_k). \end{aligned}$$

3. Gleichung (9.12) ergibt sich aus (9.8) und (9.11):

$$\begin{aligned} P(B_i|A) &= \frac{P(B_i \cap A)}{P(A)} \\ &= \frac{P(B_i) \cdot P(A|B_i)}{\sum_{k=1}^n P(B_k) \cdot P(A|B_k)}. \end{aligned}$$

□

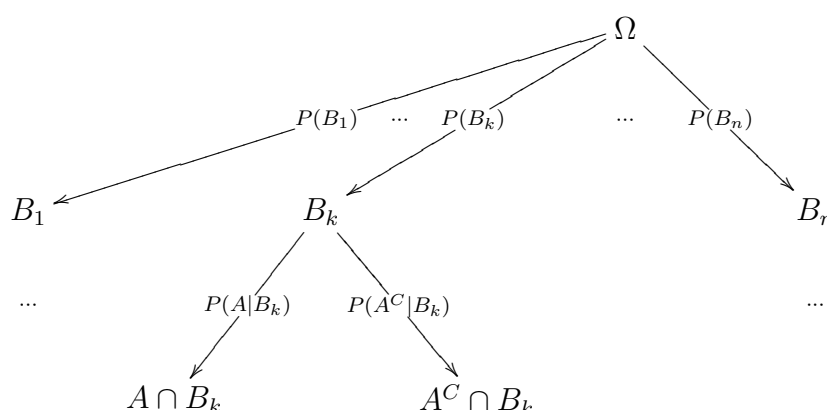


Abbildung 9.1: Wahrscheinlichkeitsbaum zur Formel der totalen Wahrscheinlichkeit

***Bemerkung 9.1.14** (Wahrscheinlichkeitsbaum).

Abbildung 9.1 illustriert Formel (9.11) der totalen Wahrscheinlichkeit mit Hilfe eines **Wahrscheinlichkeitsbaums**, der hier nur unvollständig dargestellt ist. Einige nicht eingezeichnete Kanten und Knoten (s.u.) werden durch Pünktchen angedeutet.

Der Wahrscheinlichkeitsbaum ist ein azyklischer **gerichteter Graph**, dessen **Knoten** Ereignissen entsprechen und deren **orientierte Kanten** mit Wahrscheinlichkeiten gewichtet sind: Dabei gehen von einem Knoten, z.B. dem, der dem Ereignis B_k entspricht, Kanten zu Knoten, die paarweise disjunkten Ereignissen, im Beispiel $A \cap B_k$ und $A^C \cap B_k$. Diese Kanten sind mit den bedingten Wahrscheinlichkeiten $P(A|B_k)$ und $P(A^C|B_k) = 1 - P(A|B_k)$, respektive, gewichtet.

Vom oberen Knoten, der **Wurzel**, die dem sicheren Ereignis Ω entspricht, gehen n Kanten aus, deren Zielknoten jeweils einem der Ereignisse $B_1, \dots, B_n$ entsprechen. Da genau eines dieser Ereignisse eintritt, können wir das Eintreten von B_k als eindeutig festgelegten Pfad („Spaziergang“ entlang der Kante) zum entsprechenden Knoten vorstellen. Da dies mit der Wahrscheinlichkeit $P(B_k)$ geschieht, gewichten wir die entsprechende Kante mit dieser Wahrscheinlichkeit. An dem B_k entsprechenden Knoten haben wir also die Information, dass das Ereignis B_k eintritt. Jetzt unterscheiden wir zusätzlich zwischen dem Eintreten des Ereignisses A und dessen Nicht-Eintreten, also A^C , und stellen dies in unserem Graphen durch zwei von dem B_k entsprechenden Knoten

ausgehenden Kanten mit Zielknoten $A \cap B_k$ bzw. $A^C \cap B_k$ mit den entsprechenden Gewichten $P(A|B_k)$ und $P(A^C|B_k) = 1 - P(A|B_k)$ dar. Um z.B. die Wahrscheinlichkeit $P(A \cap B_k)$ zu berechnen, gehen wir in dem Baum von der Wurzel aus den Pfad bis zum Knoten, der $A \cap B_k$ entspricht, entlang der orientierten Kanten, und multiplizieren die Gewichte der Kanten. Dadurch erhalten wir eine Formel analog zu (9.9). Wir betrachten keine weiteren Ereignisse, und somit hat unser Baum keine von den Knoten, die einem $A \cap B_k$ entsprechen, ausgehenden Kanten. Diese Knoten nennen wir **Blätter**. Wir bemerken, dass z.B. die den Ereignissen B_1 oder B_n entsprechenden Knoten *keine* Blätter sind. Wir haben nur aus Platzgründen nicht die von ihnen ausgehenden Kanten eingezeichnet.

Um nun die Wahrscheinlichkeit $P(A)$ zu berechnen, betrachten wir alle mit den Kantenorientierungen verträglichen Pfade von der Wurzel zu je einem der Blätter, die dem Eintreten von A entsprechen (also Knoten, die einem der $A \cap B_k$ entsprechen) und summieren über alle solchen Pfade die Produkte der Kantengewichte. Wir erhalten Formel (9.11).

Die gerade beschriebene Vorgehensweise kann man sich etwa wie folgt merken:

Berechnung von Wahrscheinlichkeiten mit Hilfe eines Baumdiagramms:

Multipliziere für jeden Pfad die Wahrscheinlichkeiten entlang der Kanten und summiere über alle mit dem betrachteten Ereignis verträglichen Pfade.

Bemerkung 9.1.15 (Interpretation der Formel von Bayes).

Wie durch das weiter unten folgende Beispiel 9.1.16 illustriert wird, werden in der Formel (9.12) von Bayes, die Ereignisse B_k als mögliche „Ursachen“ für das beobachtete Ereignis („Symptom“) A aufgefasst. Für jedes Ereignis B_k wird die A-priori-Wahrscheinlichkeit $P(B_k)$ als bekannt vorausgesetzt und ebenso die bedingten Wahrscheinlichkeiten dafür, dass bei Eintreten von Ursache B_k auch das Symptom A eintritt.

Mit Hilfe der Formel von Bayes wird für ein B_i die A-posteriori-Wahrscheinlichkeit berechnet unter der zusätzlichen Information, dass das Symptom A beobachtet wird.

Diese Vorgehensweise der Korrektur von A-priori-Wahrscheinlichkeiten aufgrund von Beobachtungen spielt in der *Bayesischen Statistik* eine wichtige Rolle.

Beispiel 9.1.16 (Diagnostischer Test, vgl. [Kre02]).

Eine Krankheit komme bei etwa 0,5% der Bevölkerung vor. Ein Test zur Auffindung der Krankheit führe bei 99% der Kranken zu einer Reaktion, aber auch bei 2% der Gesunden. Wir möchten die Wahrscheinlichkeit dafür ermitteln, dass eine Person, bei der die Reaktion eintritt, die Krankheit tatsächlich hat, und des Weiteren die Wahrscheinlichkeit, dass eine Person, bei der keine Reaktion eintritt, in Wirklichkeit krank ist. Dazu definieren wir mögliche Ereignisse:

$$\begin{aligned} B_1 &: \text{„Die Person hat die Krankheit.“,} \\ B_2 = B_1^C &: \text{„Die Person hat die Krankheit nicht.“,} \\ A_1 &: \text{„Test positiv“,} \\ A_2 = A_1^C &: \text{„Test negativ“}. \end{aligned}$$

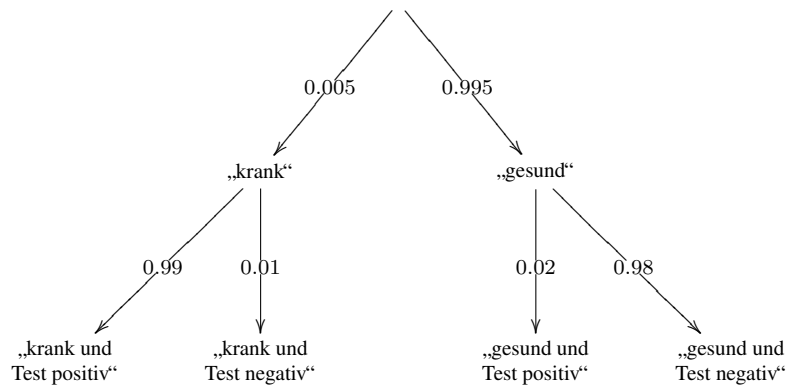


Abbildung 9.2: Wahrscheinlichkeitsbaum zum diagnostischen Test

Nach der Formel von Bayes gilt

$$\begin{aligned}
 P(B_1|A_1) &= \frac{P(B_1) \cdot P(A_1|B_1)}{P(B_1) \cdot P(A_1|B_1) + P(B_2) \cdot P(A_1|B_2)} \\
 &= \frac{5 \cdot 10^{-3} \cdot 0.99}{5 \cdot 10^{-3} \cdot 0.99 + (1 - 5 \cdot 10^{-3}) \cdot 0.02} \\
 &\approx 0.2.
 \end{aligned}$$

Die gesuchte bedingte Wahrscheinlichkeit für eine tatsächliche Erkrankung einer Person, bei der der Test positiv ist, beträgt etwa 0.2.

Auch die Wahrscheinlichkeit dafür, dass eine negativ getestete Person tatsächlich krank ist, berechnen wir nach der Formel von Bayes:

$$\begin{aligned}
 P(B_1|A_2) &= \frac{P(B_1) \cdot P(A_2|B_1)}{P(B_1) \cdot P(A_2|B_1) + P(B_2) \cdot P(A_2|B_2)} \\
 &= \frac{5 \cdot 10^{-3} \cdot 0.01}{5 \cdot 10^{-3} \cdot 0.01 + (1 - 5 \cdot 10^{-3}) \cdot 0.98} \\
 &\approx 5.1 \cdot 10^{-5}.
 \end{aligned}$$

***Definition 9.1.17** (Effizienz diagnostischer Tests, s. [Sac02]).

Wir betrachten wie in Beispiel 9.1.16 einen diagnostischen Test für eine Krankheit. Der getestete Patient kann gesund (Ereignis K^C) oder tatsächlich krank sein (Ereignis K). Der Test kann positiv ausfallen, d.h. der Patient wird als krank getestet (Ereignis T_+), oder negativ (Ereignis $T_- = T_+^C$).

1. Die **Spezifität** des Tests ist die bedingte Wahrscheinlichkeit $P(T_-|K^C)$ für einen negativen Test, wenn der Patient gesund ist.

2. Die **Sensitivität** des Tests ist die bedingte Wahrscheinlichkeit $P(T_+|K)$ für einen positiven Test, wenn der Patient krank ist.

Spezifizität und Sensitivität können wir als Gütekriterium eines Tests ansehen. Sie sollten beide nahe bei 1 liegen. Die bedingte Wahrscheinlichkeit $P(K|T_+)$ ist der **Voraussagewert** eines positiven Testergebnisses bei Kranken, und $P(K^C|T_-)$ ist der Voraussagewert eines negativen Testergebnisses bei Gesunden. Diese sollten idealerweise ebenfalls nahe bei 1 liegen. Sie hängen nach der Formel von Bayes (9.12) allerdings auch von der A-priori-Wahrscheinlichkeit für die Krankheit ab, welche als die relative Häufigkeit „Anzahl der Kranken geteilt durch die Gesamtzahl der Menschen“ (z.B. in einem bestimmten Land) definiert ist, der so genannten **Prävalenz** der Krankheit. Diese Abhängigkeit kann wie in Beispiel 9.1.16 zu niedrigen Voraussagewerten führen, wenn die Krankheit nur sehr selten ist, also zu typischem „Fehlalarm bei seltenen Ereignissen“.

9.1.3 Unabhängigkeit von Ereignissen

Beispiel 9.1.18 (für zwei unabhängige Ereignisse).

Wir betrachten folgendes Experiment: Es wird zweimal mit einem Laplace-Würfel gewürfelt. Wir betrachten das Ereignis A , dass die „Summe der Augenzahlen gerade“ und Ereignis B , dass der „zweite Wurf eine 1“ ist. Es gilt $P(A) = \frac{1}{2}$, $P(B) = \frac{1}{6}$, $P(A \cap B) = \frac{1}{12}$, wie man durch Abzählen der jeweiligen Mengen sieht. Also

$$\begin{aligned} P(A \cap B) &= P(A) \cdot P(B) \\ \Leftrightarrow P(A) &= P(A|B) \\ \Leftrightarrow P(B) &= P(B|A). \end{aligned}$$

D.h. durch die zusätzlichen Informationen, dass B eintritt, ändert sich nichts an der (bedingten) Wahrscheinlichkeit dafür, dass A eintritt.

Definition 9.1.19 (Unabhängigkeit zweier Ereignisse).

Zwei Ereignisse A und B heißen **voneinander unabhängig**, wenn die **Produktformel**

$$P(A \cap B) = P(A) \cdot P(B)$$

gilt.

***Bemerkung 9.1.20** (zum Begriff Unabhängigkeit).

1. Die Relation „ A ist unabhängig von B “ ist *symmetrisch*, d.h. „ A ist unabhängig von B “ genau dann, wenn „ B unabhängig von A “ ist. Aber im allgemeinen ist sie nicht *reflexiv* (für $0 < P(A) < 1$ gilt z.B. , dass $P(A \cap A) = P(A) \neq P(A) \cdot P(A)$) oder *transitiv* (aus „ A ist unabhängig von B “ und „ B ist unabhängig von C “ folgt i.a. *nicht*, dass „ A unabhängig von C “ ist, wie man für die Wahl eines Beispiels mit $A = C$ mit $0 < P(A) < 1$ und $B = \emptyset$ sieht.)

2. Ebenso ist die Nicht-Unabhängigkeit zweier Ereignisse nicht transitiv. Als Gegenbeispiel betrachten wir den Laplaceschen Wahrscheinlichkeitsraum (vgl. Definition 9.1.7), bestehend aus $\Omega := \{1, 2, 3, 4\}$ und der Verteilung $P(\{\omega\}) = \frac{1}{4}$ für jedes $\omega \in \Omega$ sowie die Ereignisse $A := \{1, 2\}$, $B := \{1\}$ und $C := \{1, 3\}$. Man rechnet leicht nach, dass A nicht unabhängig von B und B nicht unabhängig von C ist. Allerdings ist A unabhängig von C .
3. Die Unabhängigkeit ist als *wahrscheinlichkeitstheoretische* Unabhängigkeit zu verstehen. Durch die Information über B kann man keine bessere „Voraussage“ über A machen. In Beispiel 9.1.18 bestimmt das Ergebnis B , welches eine Aussage über den zweiten Wurf macht, in welcher Weise A eintreten kann, d.h. welche Elementarereignisse eintreten können, die Teilmengen von A sind. Bei einem nicht-fairen Würfel mit

$$\tilde{P}(\omega) = \begin{cases} \frac{1}{9} & \text{für gerade } \omega, \\ \frac{2}{9} & \text{für ungerade } \omega, \end{cases}$$

wären A und B voneinander abhängig. Es gilt dann nämlich:

$$\begin{aligned} P(A) &= \left(\frac{1}{3}\right)^2 + \left(\frac{2}{3}\right)^2 = \frac{5}{9}, \\ P(B) &= \frac{2}{3}, \\ P(A \cap B) &= \underbrace{P(B)}_{=\frac{2}{3}} \cdot \underbrace{P(\text{„erster Wurf ungerade“})}_{=\frac{2}{3}} = \frac{4}{9}, \end{aligned}$$

aber

$$P(A) \cdot P(B) = \frac{10}{27} \neq \frac{4}{9} = P(A \cap B).$$

***Definition 9.1.21. (Unabhängigkeit einer Familie von Ereignissen)**

Sei $\{A_i, i \in J\}$ eine endliche Familie von Ereignissen.

1. Wir sagen, dass die **Produktformel** für $\{A_i, i \in J\}$ gilt, wenn

$$P\left(\bigcap_{i \in J} A_i\right) = \prod_{i \in J} P(A_i). \quad (9.13)$$

2. Wir sagen, dass eine (nicht unbedingt endliche) Familie $\mathcal{A} = \{A_i, i \in I\}$ von Ereignissen **unabhängig** ist, wenn für jede endliche Teilfamilie $\{A_i, i \in J\}$ mit $J \subset I$ die Produktformel gilt.

9.1.4 Produktexperimente

Die Definitionen und Sätze in diesem Abschnitt sind recht theoretisch für diese Vorlesung und nur der Vollständigkeit halber für besonders Interessierte aufgeschrieben. Jedoch für alle wichtig ist ein gutes Verständnis der Beispiele.

Seien $(\Omega_1, P_1), \dots, (\Omega_n, P_n)$ Wahrscheinlichkeitsräume für gewisse Zufallsexperimente. Wir wollen einen Wahrscheinlichkeitsraum definieren, der die unabhängige Hintereinanderausführung dieser Experimente beschreibt.

***Definition 9.1.22** (Produkt von Wahrscheinlichkeitsräumen).

Die Menge

$$\begin{aligned}\Omega &= \prod_{i=1}^n \Omega_i = \Omega_1 \cdots \Omega_n \\ &= \{(\omega_1, \dots, \omega_n) \mid \omega_i \in \Omega_i \text{ für } i = 1, \dots, n\}\end{aligned}\tag{9.14}$$

heißt das **(kartesische) Produkt** oder auch die **Produktmenge** von $(\Omega_i)_{1 \leq i \leq n}$. Durch die Wahrscheinlichkeitsfunktion

$$P(\omega) = \prod_{i=1}^n P_i(\omega_i)\tag{9.15}$$

ist ein Wahrscheinlichkeitsmaß auf Ω definiert, das wir ebenfalls mit P bezeichnen. Wir nennen (Ω, P) das **Produkt der Wahrscheinlichkeitsräume** $(\Omega_i, P_i)_{1 \leq i \leq n}$.

***Satz 9.1.23. (Eindeutigkeit des Produkts von Wahrscheinlichkeitsräumen)**

1. Durch (9.15) ist tatsächlich ein Wahrscheinlichkeitsmaß auf Ω definiert.
2. Sei X_i die i -te Koordinatenfunktion auf Ω , d.h. $X_i(\omega) = \omega_i$. Dann gilt für $A_i \in \Omega_i$ ($i = 1, \dots, n$):

$$P\left(\bigcap_{i=1}^n \{X_i \in A_i\}\right) = \prod_{i=1}^n P_i(A_i).\tag{9.16}$$

Hierbei haben wir folgende nützliche Notation für als Urbild definierte Mengen verwendet:

$$\{X_i \in A_i\} = \{\omega = (\omega_1, \dots, \omega_n) \in \Omega \mid X_i(\omega) = \omega_i \in A_i\}.$$

Insbesondere gilt dann

$$P(\{X_n \in A_k\}) = P_k(A_k) \text{ für alle } 1 \leq k \leq n.\tag{9.17}$$

3. Das durch (9.15) definierte Wahrscheinlichkeitsmaß ist das einzige Maß auf Ω , bezüglich dessen jede Mengenfamilie $(\{X_i \in A_i\})_{1 \leq i \leq n}$ unabhängig ist und für die (9.17) gilt.

Beweis: Wir beweisen nur (9.16).

$$\begin{aligned}
 P\left(\bigcap_{i=1}^n \{X_i \in A_i\}\right) &= \sum_{\omega \in A_1 \times \dots \times A_n} \\
 &= \sum_{\omega_1 \in A_1} \dots \sum_{\omega_n \in A_n} P_1(\omega_1) \dots P_n(\omega_n) \\
 &= \left(\sum_{\omega_1 \in A_1} P_1(\omega_1)\right) \dots \left(\sum_{\omega_n \in A_n} P_n(\omega_n)\right) \\
 &= \prod_{i=1}^n P_i(A_i).
 \end{aligned}$$

□

Beispiel 9.1.24 (n-facher Münzwurf).

Wir betrachten eine Folge von n unabhängigen Einzelexperimenten, die jeweils durch die Ergebnismenge $\Omega_i = \{K, Z\}$ und das Wahrscheinlichkeitsmaß

$$P_i(\omega_i) = \begin{cases} p & \text{für } \omega_i = K, \\ 1 - p & \text{für } \omega_i = Z, \end{cases}$$

(mit $1 \leq i \leq n$) beschrieben sind. Hierbei ist $0 \leq p \leq 1$.

Die Produktmenge ist

$$\Omega = \{0, 1\}^n = \{(\omega_1, \dots, \omega_n) \mid \omega_i \in \{K, Z\}, 1 \leq i \leq n\},$$

und das Wahrscheinlichkeitsmaß ist gegeben durch seine Wahrscheinlichkeitsfunktion

$$\begin{aligned}
 P(\omega) &= \prod_{i=1}^n P_i(\omega_i) \\
 &= p^k (1 - p)^{n-k},
 \end{aligned} \tag{9.18}$$

wobei k die Anzahl der Indizes i mit $\omega_i = 1$ ist.

Definition 9.1.25 (Bernoulli-Verteilung).

Der in Beispiel 9.1.24 betrachtete Produktraum (Ω, P) heißt **Bernoulli-Experiment mit Erfolgswahrscheinlichkeit p** , und P heißt **Bernoulli-Verteilung**.

Beispiel 9.1.26 (Binomialverteilung).

Wir führen Beispiel 9.1.24 fort. Sei für $0 \leq k \leq n$ mit E_k das Ereignis bezeichnet, dass genau k -mal ein Erfolg (eine 1) eintritt. Es gibt genau $\binom{n}{k}$ solcher $\omega \in \Omega$. Also

$$P(E_k) = \binom{n}{k} p^k (1 - p)^{n-k} =: b_{n,p}(k). \tag{9.19}$$

Wir überprüfen durch eine kurze Rechnung, dass die Summe der $P(E_k)$ gleich 1 ist:

$$\begin{aligned} \sum_{k=0}^n b_{n,p}(k) &= \sum_{k=0}^n \binom{n}{k} p^k (1-p)^{n-k} \\ &= (p + (1-p))^n \\ &= 1. \end{aligned}$$

Dabei haben wir im ersten Schritt die binomische Formel verwendet.

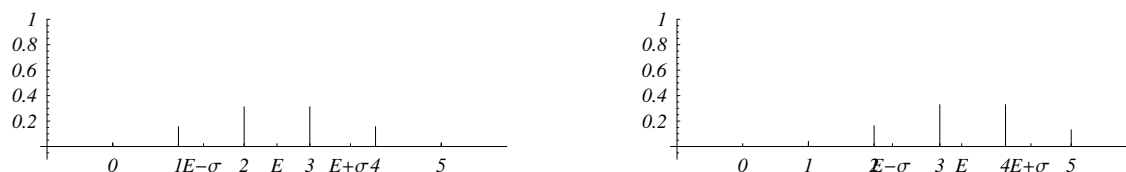


Abbildung 9.3: Stabdiagramme für die Binomialverteilungen $b_{5, \frac{1}{2}}$ und $b_{5, \frac{2}{3}}$.

Definition 9.1.27 (Binomialverteilung mit Parametern n und p).

Die durch die Zahlen $b_{n,k}(k)$ (s. (9.19)) gegebene Wahrscheinlichkeitsverteilung auf $\{0, \dots, n\}$ heißt **Binomialverteilung mit Parametern n und p** .

Beispiel 9.1.28 („Mensch ärgere Dich nicht“).

Wie groß ist die Wahrscheinlichkeit, dass bei dreimaligem Würfeln mit einem fairen Würfel keine 6 vorkommt? Wir wählen für den Wahrscheinlichkeitsraum für den i -ten Wurf

$$\Omega_i := \{\{1, 2, 3, 4, 5\}, \{6\}\}.$$

Dann gilt nach Voraussetzung (fairer Würfel):

$$P_i(\{6\}) = \frac{1}{6} = p.$$

Das Ereignis „keine 6“ entspricht der Menge

$$E_0 = \{(\omega_1, \omega_2, \omega_3) \mid \omega_i \in \{1, 2, 3, 4, 5\} \text{ für } 1 \leq i \leq 3\}.$$

Es gilt nach (9.19), dass

$$\begin{aligned} P(E_0) &= \binom{3}{1} \left(\frac{1}{6}\right)^0 \left(1 - \frac{1}{6}\right)^{3-0} \\ &= 1 \cdot 1 \cdot \left(\frac{5}{6}\right)^3 \\ &= \frac{125}{216}. \end{aligned}$$

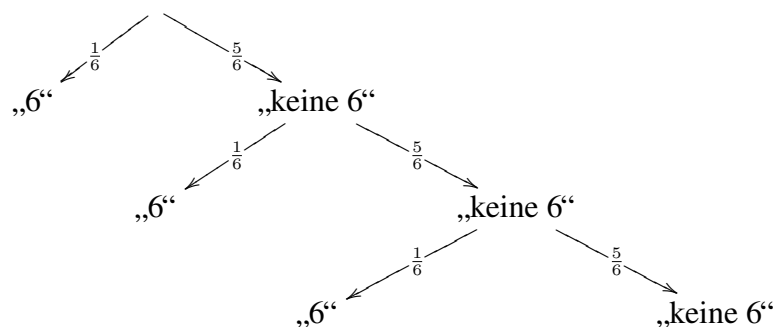


Abbildung 9.4: Graph für ein Bernoulli-Experiment

Auch in diesem Beispiel ist es hilfreich, sich die Ereignisse und Wahrscheinlichkeiten mit Hilfe eines Graphen, s. Abbildung 9.4 zu veranschaulichen. (Vgl. Bemerkung 9.1.14 sowie Abbildung 9.1.) Die Zielknoten von Kanten beschriften wir nun aber mit dem Ausgang des jeweils letzten (also dem der Kante entsprechendem) Wurf anstatt mit der gesamten Folge von bis dahin geschehenen Würfelausgängen. Zur Berechnung der Wahrscheinlichkeit eines Elementarereignisses geht man entlang dem Pfad, der zum Elementarereignis führt (dies entspricht dem Produkt von Ergebnissen einzelner Experimente (Würfe)) und multipliziert die Wahrscheinlichkeitswerte der Kanten. Alle anderen Pfade verfolgen wir daher nur bis „zur ersten 6“.

Das Produkt entlang dem Pfad, der dem Ereignis „keine 6“ entspricht, ist

$$\frac{5}{6} \cdot \frac{5}{6} \cdot \frac{5}{6} = \frac{125}{216} \quad (9.20)$$

und gleich dem oben schon berechneten Wert.

9.1.5 Zufallsvariablen

Definition 9.1.29 (Zufallsvariable).

Seien (Ω, P) ein endlicher Wahrscheinlichkeitsraum und χ eine Menge. Eine Funktion $X : \Omega \rightarrow \chi$ heißt **Zufallsexperiment mit Werten in χ** (oder auch **χ -wertige Zufallsvariable**). Falls $\chi = \mathbb{R}$, heißt X **reelle Zufallsvariable**.

Bemerkung 9.1.30 (zum Begriff „Zufallsvariable“).

Üblicherweise wird eine so genannte Unbestimmte, z.B. das Argument einer Funktion, als Variable bezeichnet. Man beachte, dass mit Zufallsvariable selber eine Funktion gemeint ist (deren Wert mit dem zufälligen Argument variiert).

Beispiel 9.1.31 (für reelle Zufallsvariablen).

1. **Geldwette bei Münzwurf:** Ein einfacher Münzwurf sei durch $\Omega = \{K, Z\}$, $P(K) = p$, $P(Z) = 1 - p$ modelliert, wobei $0 \leq p \leq 1$. Bei Kopf erhält man 2 Euro Gewinn, bei

Zahl verliert man 1 Euro. Der Gewinn (Verlust) ist eine reelle Zufallsvariable:

$$\begin{aligned} X : \Omega &\rightarrow \{-1, 2\} \in \mathbb{R}, \\ X(K) &= 2, \\ X(Z) &= -1. \end{aligned}$$

2. **Würfeln:** $\Omega = \{1, \dots, 6\}$, wobei mit $\omega = 1$ das Elementarereignis „Es wird eine 1 gewürfelt.“ gemeint ist. Sei X die Zufallsvariable, die jedem Wurf die erzielte Augenzahl zuordnet, also z.B.

$$X(1) = 1,$$

wobei die 1 auf der linken Seite das Elementarereignis „Es wird eine 1 gewürfelt.“ bezeichnet und die 1 auf der rechten Seite die reelle Zahl 1.

3. Vergleiche Beispiel 9.1.26: Wir betrachten die **Binomialverteilung** zum n -maligen Münzwurf mit Ergebnissen eines einzelnen Münzwurfes in $\{K, Z\}$. Die Anzahl der Erfolge (Kopf) sei mit $X(\omega)$ bezeichnet, also

$$\begin{aligned} X : \Omega = \{K, Z\}^n &\rightarrow \{0, \dots, n\}, \\ (\omega_1, \dots, \omega_n) &\mapsto \sum_{i=1}^n X_i(\omega), \end{aligned} \tag{9.21}$$

wobei

$$\begin{aligned} X : \Omega &\rightarrow \{0, 1\}, \\ X_i(\omega) &= \begin{cases} 1 & \text{für } w_i = K, \\ 0 & \text{für } w_i = Z. \end{cases} \end{aligned}$$

Die Zufallsvariable X ist also die Summe der Zufallsvariablen X_i .

Satz 9.1.32. (Eine Zufallsvariable definiert eine Wahrscheinlichkeitsfunktion auf dem Bildraum)

Seien (Ω, P) ein endlicher Wahrscheinlichkeitsraum und $X : \Omega \rightarrow \chi$ eine Zufallsvariable. Dann ist auf χ eine Wahrscheinlichkeitsfunktion P_X durch

$$\begin{aligned} P_X : \chi &\rightarrow [0, 1], \\ P_X(y) &= P(\{X = y\}) \\ &= \sum_{\omega \in \Omega, X(\omega)=y} P(\omega) \end{aligned}$$

definiert. Hierbei bezeichnet $\{X = y\} := \{\omega \in \Omega | X(\omega) = y\}$ die Urbildmenge von y bezüglich der Abbildung X .

Definition 9.1.33 (Verteilung einer Zufallsvariablen).

Das Wahrscheinlichkeitsmaß zur Wahrscheinlichkeitsfunktion P_X aus Satz 9.1.32 heißt **Verteilung von X bezüglich P** oder auch das **Wahrscheinlichkeitsmaß von X bezüglich P** .

Bemerkung 9.1.34 (Wichtigkeit von Verteilungen).

Meistens interessiert man sich ausschließlich für die Verteilung von Zufallsvariablen X und nicht für das Wahrscheinlichkeitsmaß P auf Ω . Wir hatten schon in Beispiel 9.1.8 gesehen, dass verschiedene Wahlen von Ω möglich sein können. Oftmals ist der „steuernde Wahrscheinlichkeitsraum“ nicht explizit bekannt oder sehr kompliziert.

Beispiel 9.1.35 (Binomialverteilung als Verteilungsmaß).

Das in (9.19) durch die Binomialverteilung definierte Wahrscheinlichkeitsmaß P auf der Menge $\{E_0, \dots, E_n\}$ können wir offensichtlich auch als die Verteilung der Zufallsvariablen X aus (9.21) in Beispiel 9.1.31.3 auffassen, also als Wahrscheinlichkeitsmaß auf der Menge $\{0, 1, \dots, n\}$. Ein Element k aus dieser Menge entspricht dabei der Menge E_k aus Beispiel 9.1.31.3. Also

$$P_X(k) = b_{n,p}(k).$$

***Definition 9.1.36** (Unabhängigkeit von Zufallsvariablen).

Sei (Ω, P) ein endlicher Wahrscheinlichkeitsraum. Eine Familie $(X_i)_{i \in I}$ von Zufallsvariablen $X_i : \Omega \rightarrow \chi_i$ (mit $i \in I$) heißt **unabhängig**, wenn für jede endliche Teilmenge $J \subset I$ und jede Wahl von $A_j \subset \chi_j$ für alle $j \in J$ die Familie $(\{X_j \in A_j\})_{j \in J}$ unabhängig ist. (vgl. Definition 9.1.21).

Bemerkung 9.1.37 (Interpretation der Unabhängigkeit von Zufallsvariablen).

Im Folgenden wird uns die Unabhängigkeit von Zufallsvariablen meistens als Voraussetzung für mathematische Sätze begegnen. Die Folgerungen aus der Unabhängigkeit sind sehr nützlich und auch nicht so abstrakt wie Definition 9.1.36. Jeder sollte zumindest folgende Interpretation der Unabhängigkeit von zwei Zufallsvariablen verstehen:

Seien z.B. X_1 und X_2 zwei voneinander unabhängige Zufallsvariablen mit Werten in χ_1 und χ_2 , respektive. Die Verteilung von X_2 können wir als „Voraussage“ über den zufälligen Wert von X_2 interpretieren. (vgl. Bemerkung 9.1.3.) Seien $A_2 \subset \chi_2$ und $x_1 \in \chi_1$ mit $P(\{X_1 = x_1\}) > 0$. Die Kenntnis, dass X_1 den Wert x_1 annimmt, ermöglicht uns keine „bessere“ Voraussage über den Wert von X_2 . Dies wird an Beispiel 9.1.39 veranschaulicht werden.

***Bemerkung 9.1.38. (Produktformel für unabhängige Zufallsvariablen)**

Für unabhängige Zufallsvariablen $X_1, \dots, X_n$ mit $X_i : \Omega \rightarrow \chi_i$ gilt

$$P(X_1 \in A_1 \wedge \dots \wedge X_n \in A_n) = \prod_{i=1}^n P(X_i \in A_i)$$

für jede Wahl von Ereignissen $A_i \subset \chi_i$. Die Berechnung der Wahrscheinlichkeit von solchen Ereignissen der Form $\{X_1 \in A_1\} \cap \dots \cap \{X_n \in A_n\}$ ist also besonders einfach.

***Beispiel 9.1.39** (Voneinander unabhängige Münzwürfe).

Wir betrachten den zweifachen Münzwurf aus Beispiel 9.1.24 (also $n = 2$). Auf $\Omega = \{K, Z\}^2$ ist das Produktmaß gerade so definiert, dass die beiden Zufallsvariablen

$$\begin{aligned} X_i : \Omega &\rightarrow \{K, Z\}, \\ (\omega_1, \omega_2) &\mapsto \omega_i, \end{aligned}$$

von denen X_1 gerade den Ausgang des ersten Wurfs beschreibt und X_2 den des zweiten, voneinander unabhängig sind, was anschaulich auch klar sein sollte. Es gilt z.B.

$$\begin{aligned} P(\{X_1 = K \wedge X_2 = K\}) &= P_1(K) \cdot P_2(K) \\ &= P(\{X_1 = K\}) \cdot P(\{X_2 = K\}), \end{aligned}$$

wobei wir im ersten Schritt die Produktformel (9.18) für die Wahrscheinlichkeitsfunktion verwendet haben.

9.1.6 Erwartungswert, Varianz, Kovarianz

In einem Spiel wie in Beispiel 9.1.31.1 interessiert uns der zu erwartende Gewinn und allgemein der „mittlere Wert“ einer reellen Zufallsvariablen.

Definition 9.1.40 (Erwartungswert einer reellen Zufallsvariablen).

Sei X eine reelle Zufallsvariable auf dem Wahrscheinlichkeitsraum (Ω, P) . Der **Erwartungswert von X** ist definiert als

$$EX := E(X) := \sum_{\omega \in \Omega} X(\omega) \cdot P(\omega) \quad (9.22)$$

$$= \sum_{x \in \mathbb{R}} x \cdot P_X(x). \quad (9.23)$$

Bemerkung 9.1.41 (Erwartungswert einer Verteilung).

In (9.23) ist P_X die Verteilung von X (s. Definition 9.1.33). Lediglich solche Summanden sind ungleich 0, für die $P_X(x) > 0$. Dies sind aber nur endlich viele, da der Definitionsbereich und somit der Bildbereich von X endlich ist. In (9.23) wird der „steuernde Wahrscheinlichkeitsraum“ Ω nicht explizit erwähnt. Der Erwartungswert ist also eine Eigenschaft der Verteilung. (Vgl. hierzu Bemerkung 9.1.34.) Durch (9.23) ist der **Erwartungswert der Verteilung P_X** definiert, und analog definiert man allgemein den **Erwartungswert eines Wahrscheinlichkeitsmaßes auf endlichen Mengen reeller Zahlen**.

***Bemerkung 9.1.42.** (Erwartungswert einer vektorwertigen Zufallsvariablen)

Wir können in (9.22) die mit den Wahrscheinlichkeiten gewichtete Summe bilden, da die Werte $X(\omega)$ reelle Zahlen sind. Etwas allgemeiner kann man auch den Erwartungswert z.B. von Zufallsvariablen mit Werten in den komplexen Zahlen oder in reellen oder komplexen Vektorräumen.

Satz 9.1.43 (Eigenschaften des Erwartungswertes).

1. Der Erwartungswert ist *linear*, d.h. für reelle Zufallsvariablen X, Y und $\lambda \in \mathbb{R}$ gilt

$$E(\lambda X + Y) = \lambda \cdot E(X) + E(Y). \quad (9.24)$$

2. Sind X, Y unabhängig, so gilt

$$E(X \cdot Y) = E(X) \cdot E(Y).$$

Hierbei bezeichnet $X \cdot Y$ das Produkt der beiden Zufallsvariablen. Diese durch $(X \cdot Y)(\omega) = X(\omega) \cdot Y(\omega)$ definierte Produktfunktion ist wieder eine reelle Zufallsvariable auf demselben Wahrscheinlichkeitsraum.

Beispiel 9.1.44 (für Erwartungswerte spezieller Verteilungen).

1. Wir berechnen den Erwartungswert der Zufallsvariablen X aus Beispiel 9.1.31.1, also den „zu erwartenden **Gewinn beim Münzwurf**“:

$$\begin{aligned} E(X) &= p \cdot 2 + (1 - p) \cdot (-1) \\ &= -1 + 2p. \end{aligned}$$

2. Wir berechnen den Erwartungswert der **Binomialverteilung** zu den Parametern n und p (s. 9.19) auf zwei verschiedene Weisen.

1. *Methode:*

$$\begin{aligned} E(X) &= \sum_{k=0}^n k \binom{n}{k} p^k (1-p)^{n-k} \\ &= np \sum_{k=1}^n \frac{(n-1)!}{(k-1)!((n-1)-(k-1))!} p^{(k-1)} (1-p)^{((n-1)-(k-1))} \\ &= np \sum_{\tilde{k}=0}^{\tilde{n}} \binom{\tilde{n}}{\tilde{k}} p^{\tilde{k}} (1-p)^{\tilde{n}-\tilde{k}} \\ &= np (p + (1-p))^{\tilde{n}} \\ &= np. \end{aligned}$$

Dabei haben wir die Substitution $n-1 = \tilde{n}$ und $k-1 = \tilde{k}$ verwendet.

2. *Methode:* Wir verwenden (9.24) (Linearität von E). Es gilt

$$X = X_1 + \cdots + X_n$$

mit $X_i : \Omega \rightarrow \{0, 1\}$, $P(\{X_i = 1\}) = p$, $P(\{X_i = 0\}) = 1 - p$, also $E(X_i) = p$ und somit

$$\begin{aligned} E(X) &= \sum_{i=1}^n E(X_i) \\ &= np. \end{aligned}$$

3. Wir berechnen den Erwartungswert für die Augenzahl beim **Laplace-Würfel**, gegeben durch $\Omega = \{1, \dots, 6\}$ und $P(\omega) = \frac{1}{6}$ für $\omega \in \Omega$. Die Zufallsvariable X gibt die Augenzahl an. (S. Beispiel 9.1.31.2.) Wir erhalten

$$E(X) = \sum_{i=1}^6 i \cdot \frac{1}{6} = 3.5. \quad (9.25)$$

Insbesondere sehen wir, dass der Erwartungswert i.a. nicht als Wert von der Zufallsvariablen angenommen wird.

4. Wir vergleichen das letzte Beispiel mit der Zufallsvariablen Y , definiert auf demselben (Ω, P) durch

$$Y(\omega) = 3.5 \quad \text{für } \omega \in \{1, \dots, 6\}.$$

Diese Zufallsvariable hat den gleichen Erwartungswert wie der Laplace-Würfel:

$$E(Y) = 3.5.$$

Dennoch sind die beiden Zufallsvariablen nicht gleichverteilt. Wie durch die **Stabdiagramme** in Abbildung 9.5 veranschaulicht wird, ist die Verteilung P_y deterministisch, wohingegen P_x um den Erwartungswert streut.

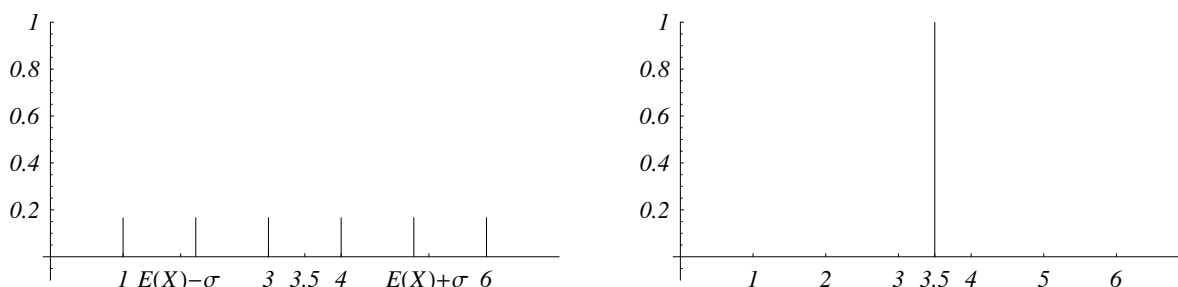


Abbildung 9.5: Stabdiagramme für den Laplace-Würfel (links) und für eine deterministische Zufallsvariable (rechts)

Wie Beispiel 9.1.44.4 zeigt, ist eine Wahrscheinlichkeitsverteilung in den reellen Zahlen nicht allein durch ihren Erwartungswert charakterisiert. Dies motiviert die Einführung von weiteren Kenngrößen von Zufallsvariablen.

Definition 9.1.45. (Varianz, Streuung, Kovarianz, Korrelationskoeffizient)

Seien (Ω, P) ein endlicher Wahrscheinlichkeitsraum und X, Y reelle Zufallsvariablen.

1. Die **Varianz** von X ist

$$\text{Var}(X) = E((X - E(X))^2).$$

2. Die **Streuung** (oder **Standardabweichung**) von X ist

$$\sigma = \sqrt{\text{Var}(X)}.$$

3. Die **Kovarianz** von X und Y ist

$$\text{Cov}(X, Y) = E((X - E(X)) \cdot (Y - E(Y))).$$

4. Der **Korrelationskoeffizient** von X und Y (mit $\sigma_x, \sigma_y \neq 0$) ist

$$\rho_{X,Y} = \frac{\text{Cov}(X, Y)}{\sigma_x \sigma_y}. \quad (9.26)$$

5. Zufallsvariablen X, Y mit $\text{Cov}(X, Y) = 0$ heißen **unkorreliert**.

Satz 9.1.46 (Eigenschaften von Varianz und Kovarianz).

Seien X, Y, X_i (für $1 \leq i \leq n$) reelle Zufallsvariablen und $a, b, c, d \in \mathbb{R}$. Dann gilt:

1.
$$\text{Var}(X) = E(X^2) - (E(X))^2. \quad (9.27)$$

2.
$$\text{Var}(aX + b) = a^2 \cdot \text{Var}(X). \quad (9.28)$$

3.
$$\text{Cov}(X, Y) = E(XY) - E(X) \cdot E(Y). \quad (9.29)$$

4.
$$\text{Cov}(aX + b, cY + d) = a \cdot c \cdot \text{Cov}(X, Y), \quad (9.30)$$

5.
$$\text{Var}(X_1 + \cdots + X_n) = \sum_{i=1}^n \text{Var}(X_i) + \sum_{\substack{(i,j), \\ i \neq j}} \text{Cov}(X_i, X_j), \quad (9.31)$$

wobei in der letzten Summe die Summanden $\text{Cov}(X_1, X_2)$ und $\text{Cov}(X_2, X_1)$ etc. auftreten.

6. Sind X, Y unabhängig, so sind sie auch unkorreliert.

7. (**Formel von Bienaymé**) Wenn $X_1, \dots, X_n$ unabhängig sind, dann gilt

$$\text{Var}(X_1 + \cdots + X_n) = \sum_{i=1}^n \text{Var}(X_i). \quad (9.32)$$

Bemerkung 9.1.47. (Aus Unkorreliertheit folgt nicht Unabhängigkeit)

Die Umkehrung von Satz 9.1.46.6 gilt nicht, d.h. aus der Unkorreliertheit von Zufallsvariablen folgt im Allgemeinen *nicht* deren Unabhängigkeit, wie wir in Beispiel 9.1.53.3 sehen werden.

Beispiel 9.1.48 (Varianz bei der Augenzahl des Laplace-Würfels).

Es gilt für das **zweite Moment** der Augenzahl X des Laplace-Würfels:

$$E(X^2) = \sum_{i=1}^6 i^2 \cdot \frac{1}{6} = \frac{91}{6}.$$

Daraus erhalten wir nach (9.27) und unter Verwendung von (9.25)

$$\begin{aligned} \text{Var}(X) &= E(X^2) - (E(X))^2 \\ &= \frac{91}{6} - 3.5^2 \\ &= \frac{35}{12}. \end{aligned} \tag{9.33}$$

Die Streuung ist also $\sigma_X \approx 1.71$.

Beispiel 9.1.49 (Varianz der Binomialverteilung).

Mit Hilfe der Formel von Bienaymé (9.32) berechnen wir analog zur 2. Methode in Beispiel 9.1.44.2 die Varianz der Binomialverteilung zu den Parametern n und p . Die Varianz von X_i ist

$$\begin{aligned} \text{Var}(X_i) &= (0 - E(X_i)) \cdot P(X_i = 0) + (1 - E(X_i)) \cdot P(X_i = 1) \\ &= (-p)^2 \cdot (1 - p) + (1 - p)^2 \cdot p \\ &= p(1 - p). \end{aligned}$$

Aus der Unabhängigkeit der X_i folgt also

$$\begin{aligned} \text{Var}(X) &= \text{Var}\left(\sum_{i=1}^n X_i\right) \\ &= \sum_{i=1}^n \text{Var}(X_i) \\ &= n p (1 - p). \end{aligned} \tag{9.34}$$

Zur Veranschaulichung von Korrelation führen wir noch den wichtigen Begriff der *gemeinsamen Verteilung* ein und beschränken uns dabei hier auf den Fall zweier reellwertiger Zufallsvariablen. Zur naheliegenden Verallgemeinerung auf den Fall von endlich vielen Zufallsvariablen mit Werten in beliebigen Mengen s. z.B. [Kre02]

Definition 9.1.50. (Gemeinsame Verteilung zweier reeller Zufallsvariablen)

Seien $X, Y : \Omega \mapsto \mathbb{R}$ zwei auf derselben Ergebnismenge Ω definierten reellwertigen Zufallsvariablen. Die Verteilung $P_{X \times Y}$ (vgl. Definition 9.1.33) der Produktfunktion

$$X \times Y : \Omega \mapsto \mathbb{R}^2$$

heißt **gemeinsame Verteilung** von X und Y . Die Funktion $X \times Y$ nimmt genau die Werte $(x, y) \in \mathbb{R}^2$ mit positiver Wahrscheinlichkeit an, für die $P_X(x) > 0$ und $P_Y(y) > 0$ gilt und gemäß Satz 9.1.32 erhalten wir

$$P_{X \times Y}(x, y) = P(\omega \in \Omega : X(\omega) = x \text{ und } Y(\omega) = y).$$

Beispiel 9.1.51 (Korrelation bei Merkmalsverteilung).

Wir betrachten ein einfaches Zahlenbeispiel für eine gemeinsame Verteilung zweier Zufallsvariablen und berechnen deren Korrelationskoeffizient. Die Zufallsvariablen nehmen hier jeweils nur zwei Werte an, die wir beliebig mit 0 und 1 gewählt haben. Solche Zufallsvariablen könnten z.B. Merkmalsausprägungen in einer Population beschreiben, wobei man nur zwischen zwei verschiedenen Stufen der Ausprägung je Individuum und Merkmal unterscheidet, nämlich „Merkmal vorhanden“ und „Merkmal nicht vorhanden“, also z.B. Linkshändigkeit (Wert 0 für Linkshänder und 1 für Rechtshänder) oder Kurzsichtigkeit (kurzsichtig oder nicht kurzsichtig). Ein Korrelationskoeffizient nahe bei 1 oder -1 deutet im Sinne von Bemerkung 9.1.52.2 auf einen linearen Zusammenhang zwischen den Merkmalen hin.

Achtung: Wir weisen ausdrücklich darauf hin, dass man in der *Statistik* keine Wahrscheinlichkeiten gegeben hat, sondern relative Häufigkeiten in einer Stichprobe, aus denen man die Wahrscheinlichkeiten schätzen kann. Solche z.B. durch Zählungen gewonnenen Daten werden demzufolge auch anders ausgewertet als hier beschrieben, insbesondere wenn die absolute Anzahl der Beobachtungen oder Experimente klein ist. Näheres dazu ist u.a. in den von uns empfohlenen Büchern über Statistik zu finden, etwa unter den Stichwörtern *Vierfeldertafel* oder allgemeiner *Kontingenztafel*.

Nun zum Zahlenbeispiel, anhand dessen wir lediglich die Rechnungen vorführen wollen, ohne jede weitere Interpretation.

Seien X_1 und X_2 Zufallsvariablen mit Werten in $\{0, 1\}$. Die Produktzufallsvariable $X_1 \times X_2$ nehme die Werte $(0, 0)$, $(1, 0)$, $(0, 1)$ und $(1, 1)$ mit den Wahrscheinlichkeiten $\frac{1}{10}$, $\frac{1}{5}$, $\frac{3}{10}$, $\frac{2}{5}$, respektive, an. Wir schreiben abkürzend $P_{X_1 \times X_2}(1, 1)$ statt $P_{X_1 \times X_2}(\{(1, 1)\})$ etc. Wir stellen die gemeinsame Verteilung sowie die Verteilungen von X_1 und X_2 tabellarisch dar:

	$X_2 = 0$	$X_2 = 1$	Verteilung von X_2 :
$X_1 = 0$	$\frac{1}{10}$	$\frac{3}{10}$	$\frac{2}{5}$
$X_1 = 1$	$\frac{1}{5}$	$\frac{2}{5}$	$\frac{3}{5}$
Verteilung von X_1 :	$\frac{3}{10}$	$\frac{7}{10}$	

Die Verteilung von X_1 und X_2 steht offensichtlich im oberen linken Teil der Tabelle. Die Verteilung von X_1 steht in der unteren Zeile. Die Werte wurden als Summe der Zahlen der jeweiligen Spalten berechnet. Ebenso steht die Verteilung von X_2 in der rechten Spalte. Diese Werte sind jeweils die Zeilensummen (aus dem Tabellenteil der gemeinsamen Verteilung). Eine Kontrollrechnung zeigt, dass die Summe der Werte der unteren Zeile (der rechten Spalte) jeweils 1 ergeben.

Wir berechnen nun die Kenngrößen der Verteilungen.

$$\begin{aligned}E(X_1) &= 0 \cdot \frac{2}{5} + 1 \cdot \frac{3}{5} \\&= \frac{3}{5}, \\E(X_1^2) &= \frac{3}{5}, \\\text{Var}(X_1) &= \frac{3}{5} - \left(\frac{3}{5}\right)^2 \\&= \frac{6}{25}, \\\sigma_{X_1} &= \sqrt{\frac{6}{25}} \\&\approx 0.49.\end{aligned}$$

$$\begin{aligned}E(X_2) &= \frac{7}{10}, \\E(X_2^2) &= \frac{7}{10}, \\\text{Var}(X_2) &= \frac{7}{10} - \left(\frac{7}{10}\right)^2 \\&= \frac{21}{100}, \\\sigma_{X_2} &= \sqrt{\frac{21}{100}} \\&\approx 0.46.\end{aligned}$$

$$\begin{aligned}E(X_1 \cdot X_2) &= \frac{2}{5}, \\\text{Cov}(X_1, X_2) &= E(X_1 \cdot X_2) - E(X_1) \cdot E(X_2) \\&= \frac{2}{5} - \frac{3}{5} \cdot \frac{7}{10} \\&= -\frac{1}{50},\end{aligned}$$

$$\begin{aligned}\rho_{X_1, X_2} &= \frac{-\frac{1}{50}}{\sqrt{\frac{6}{25} \cdot \frac{21}{100}}} \\ &\approx -0.089.\end{aligned}$$

Die Zufallsvariablen X_1 und X_2 sind nicht voneinander unabhängig, da ihre Kovarianz ungleich 0 ist. (Es gilt nämlich: „Unabhängigkeit $\Rightarrow$ Kovarianz gleich 0“.) Der Betrag ihres Korrelationskoeffizienten ist allerdings auch nicht besonders groß, d.h. nahe bei 0.

Bemerkung 9.1.52 (Interpretation von Korrelation).

1. (geometrische Sichtweise)

Wir können die Kovarianz als Skalarprodukt in $\mathbb{R}^n$ mit $n = |\Omega|$ auffassen (s. Definition 6.4.1) Hierzu nehmen wir an, dass alle Elementarereignisse eine positive Wahrscheinlichkeit haben. Dann gilt die Cauchy-Schwarz-Ungleichung (vgl. (6.8))

$$\text{Cov}(X, Y) \leq \sigma_x \sigma_y$$

und somit für $\sigma_x, \sigma_y \neq 0$:

$$-1 \leq \rho_{X,Y} \leq 1.$$

Den Korrelationskoeffizienten können wir dann als „Kosinus des nicht-orientierten Winkels zwischen X und Y “ auffassen.

2. (Korrelation als linearer Zusammenhang)

Für zwei Zufallsvariablen X und Y deutet ein Korrelationskoeffizient $\rho_{X,Y}$ nahe bei 1 auf eine „Tendenz“ der Variablen $X - E(X)$ und $Y - E(Y)$ hin, gemeinsam große bzw. kleine bzw. stark negative Werte anzunehmen, also auf einen „linearen Zusammenhang“. Analoges gilt für $\rho_{X,Y}$ nahe bei -1 . Wir veranschaulichen dies in Beispiel 9.1.53.

3. (Fehlinterpretationen von Korrelation)

In der Statistik wird die (*empirische*) *Korrelation* von durch Stichproben ermittelten Verteilungen betrachtet, um diese auf mögliche Zusammenhänge zu untersuchen. Bei der Interpretation starker Korrelationen sollte man jedoch sehr vorsichtig sein. Eine solche kann i.a. *nicht* als *kausaler Zusammenhang* zwischen zwei Größen gedeutet werden.

Ein prominentes Beispiel hierfür ist die Anzahl der Störche und der Neugeborenen pro Jahr in einem Land mit zunehmender Industrialisierung. Sinken in einem beobachteten Zeitraum diese beiden Werte, so sollte man daraus nicht folgern, dass die Neugeborenen von den Klapperstörchen gebracht würden, also die Zahl der Störche die Zahl der Neugeborenen kausal beeinflösse. Eine Erklärung der beobachteten Werte durch Änderung der Familienstruktur und der Verkleinerung der Lebensräume für Störche, bedingt durch Industrialisierung, also eine dritte Größe, welche die beiden anderen auf eine noch zu präzisierende Weise beeinflusst, erscheint hier sinnvoller.

Für weitere Diskussion und Beispiele verweisen wir auf [Kre02], [Sac02], [SR94] und [Sta02]. Als Stichwörter zum Nachschlagen in deutschsprachigen Büchern über Statistik seien hier *kausale Korrelation*, *Inhomogenitätskorrelation* und *Scheinkorrelation* genannt.

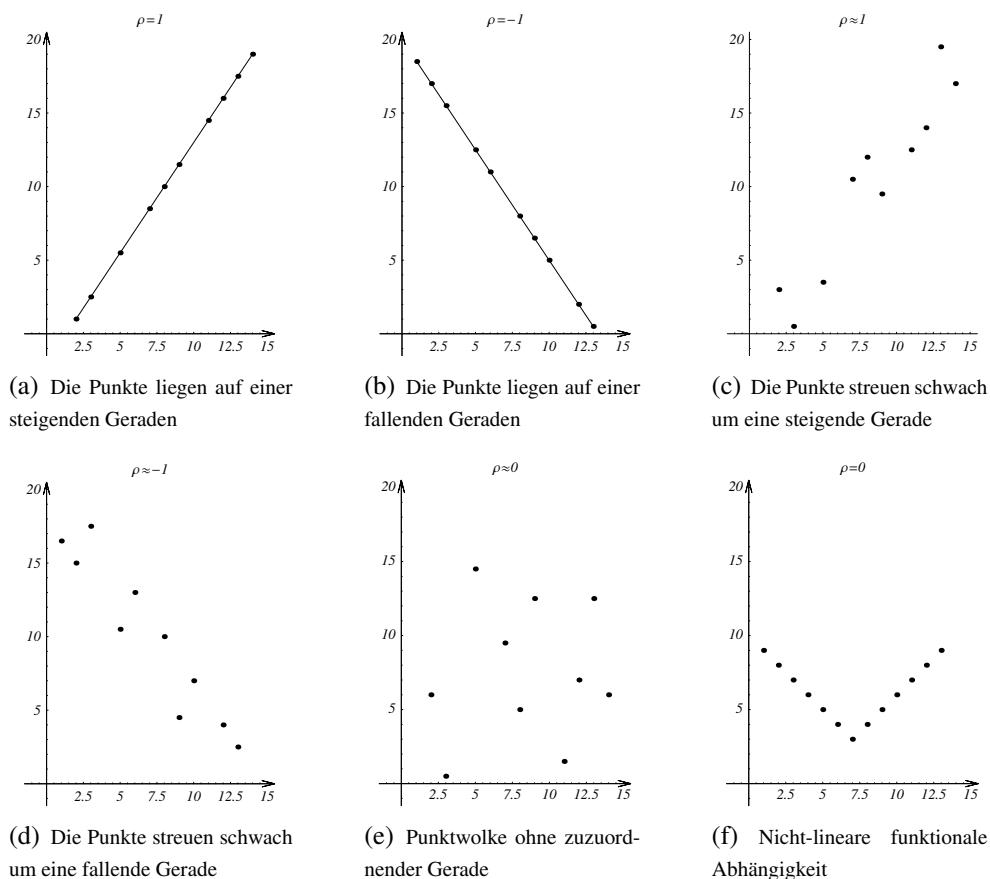


Abbildung 9.6: Illustration von Korrelationskoeffizienten mit Hilfe von gemeinsamen Verteilungen

Beispiel 9.1.53. (Illustration von speziellen gemeinsamen Verteilungen und Korrelation)

Die hier diskutierten Beispiele für gemeinsame Verteilungen sind in Abbildung 9.6 graphisch dargestellt. Die Werte der jeweiligen Verteilungen mit positiver Wahrscheinlichkeit sind als Punkte in die x - y -Ebene eingezeichnet, wobei (x, y) Werte der Funktion $X \times Y$ sind. Eine solche Darstellung könnte noch präzisiert werden, indem man zu jedem Punkt die Wahrscheinlichkeit schreibt, was bei einer kleinen Anzahl von Punkten noch übersichtlich wäre. Der Einfachheit halber habe hier jeweils jeder Punkt die gleiche Wahrscheinlichkeit.

1. Sei X eine Zufallsvariable mit Varianz $\sigma_X^2 > 0$ und sei $Y = aX + b$ mit $a \neq 0$. Wir berechnen unter Verwendung der Sätze 9.1.43 und 9.1.46 den Korrelationskoeffizienten

von X und Y .

$$\begin{aligned}\text{Var}(Y) &= a^2 \text{Var}(x), \\ \Rightarrow \sigma_Y &= |a| \cdot \sigma_X, \\ \text{Cov}(X, Y) &= \text{Cov}(X, aX + b) \\ &= a \text{Cov}(X, X) \\ &= a \sigma_X^2, \\ \rho_{X,Y} &= \frac{a \sigma_X^2}{\sigma_X |a| \sigma_X} \\ &= \text{sign}(a).\end{aligned}$$

Der Korrelationskoeffizient $\rho_{X,Y}$ ist also 1 oder -1 , je nachdem, ob a positiv oder negativ ist. Vgl. dazu auch Bemerkung 9.1.52.2. In den Abbildungen 9.6.9.6(a) und 9.6.9.6(b) sind Beispiele für solche gemeinsamen Verteilungen von X und Y dargestellt. Die Punkte der gemeinsamen Verteilung liegen auf einer Geraden. Wir bemerken auch, dass im Fall $a = 0$, also $Y = b$, die Zufallsvariable Y deterministisch ist und somit Varianz Null hat. Auch hier liegen die Punkte der gemeinsamen Verteilung von X und Y auf einer Geraden (nicht abgebildet), aber der Korrelationskoeffizient ist im Sinne von Definition 9.1.45.4 nicht definiert.

2. In den Abbildungen 9.6.9.6(c) und 9.6.9.6(d) sind die gemeinsamen Verteilungen von Zufallsvariablen dargestellt, deren Korrelationskoeffizient nahe bei 1 bzw. nahe bei -1 liegt. Die Punkte liegen zwar nicht auf einer Geraden, aber man kann jeder der Verteilungen eine Gerade zuordnen, von der die Punkte „nicht allzu sehr“ abweichen. Eine solche Zuordnung geschieht z.B. mit Hilfe von *linearer Regression*.
3. Der in Abbildung 9.6.9.6(e) dargestellten Verteilung wäre optisch nur schwer eine Gerade zuzuordnen. Der Korrelationskoeffizient in diesem Beispiel liegt nahe bei 0.
4. Wir betrachten nun noch ein sehr spezielles Beispiel. Die gemeinsame Verteilung von X und Y sei

$$P_{X \times Y}(-1, 1) = P_{X \times Y}(0, 0) = P_{X \times Y}(1, 1) = \frac{1}{3}$$

dargestellt. Die Kovarianz von X und Y ist

$$\begin{aligned}\text{Cov}(X, Y) &= \sum_{(x,y)} x \cdot y \cdot P_{X \times Y}(x, y) \\ &= \frac{1}{3} \cdot (1 \cdot (-1) + 0 \cdot 0 + 1 \cdot 1) \\ &= 0.\end{aligned}$$

Dabei haben wir in der ersten Zeile über alle Werte (x, y) mit positiver Wahrscheinlichkeit summiert. Die beiden Zufallsvariablen sind also nicht korreliert. Ihr Korrelationskoeffizient ist gleich 0.

Wir bemerken noch, dass Y *nicht* unabhängig von X ist (s. Definition 9.1.36). Im Gegenteil, es besteht sogar ein funktionaler Zusammenhang zwischen beiden Variablen. Kennt man den Wert von X , so auch den von Y . Dieser Zusammenhang ist aber *nicht* linear (vgl. 9.1.52).

Analog zu diesem Beispiel sind die Zufallsvariablen, deren gemeinsame Verteilung in Abbildung 9.6.9.6(f) dargestellt ist, unkorreliert, obwohl ein funktionaler Zusammenhang zwischen ihnen besteht.

9.1.7 Das schwache Gesetz der großen Zahlen

In diesem Abschnitt formulieren wir mit Satz 9.1.55 eine Version des schwachen Gesetzes der großen Zahlen, das insbesondere einen Zusammenhang zwischen dem abstrakt eingeführten Begriff der Wahrscheinlichkeit und relativen Häufigkeiten bei einer Folge aus lauter voneinander unabhängigen Zufallsexperimenten herstellt, die alle den gleichen Erwartungswert haben.

Der folgende Satz liefert uns eine Abschätzung für die Wahrscheinlichkeit der Abweichung einer Zufallsvariablen von ihrem Erwartungswert um mehr als eine vorgegebene Konstante. Diese Abschätzung benutzt nur die Varianz der Zufallsvariablen, ohne irgendwelche weiteren Bedingungen an die Verteilung zu stellen, und ist damit anwendbar sobald man die Varianz kennt. Allerdings ist sie in vielen Fällen auch nur sehr grob oder gar völlig nutzlos, z.B. wenn die rechte Seite in (9.35) größer gleich 1 ist. Dennoch liefert sie uns einen sehr einfachen Beweis des schwachen Gesetzes der großen Zahlen.

Satz 9.1.54 (Tschebyscheff-Ungleichung).

Sei X eine reelle Zufallsvariable auf (Ω, P) . Dann gilt für jedes $\epsilon > 0$:

$$P(|X - E(X)| > \epsilon) \leq \frac{\text{Var}(X)}{\epsilon^2}. \quad (9.35)$$

Beweis: Sei $Z = X - E(X)$. Wir definieren zu Z^2 eine Minorante, d.h. eine Zufallsvariable Y mit $Y(\omega) \leq (Z(\omega))^2$:

$$Y(\omega) := \begin{cases} 0 & \text{für } |Z(\omega)| < \epsilon, \\ \epsilon^2 & \text{für } |Z(\omega)| \geq \epsilon. \end{cases}$$

Mit Hilfe dieser Minorante können wir den Erwartungswert von Z^2 nach unten abschätzen:

$$\begin{aligned} \text{Var}(X) &= E(Z^2) \\ &\geq E(Y) \\ &= \epsilon^2 \cdot P(Y = \epsilon^2) \\ &= \epsilon^2 \cdot P(|X - E(x)| \geq \epsilon). \end{aligned}$$

□

Satz 9.1.55 (Das schwache Gesetz der großen Zahlen).

Seien $X_1, X_2, \dots$ unabhängige Zufallsvariablen mit den gleichen Erwartungswerten $E(X_1)$ und $\text{Var}(X_i) \leq M$. Dann gilt

$$P\left(\left|\frac{1}{n}(X_1 + \dots + X_n) - E(X_1)\right| \geq \epsilon\right) \leq \frac{M}{\epsilon^2 n}, \quad (9.36)$$

insbesondere

$$\lim_{n \rightarrow \infty} P\left(\left|\frac{1}{n}(X_1 + \dots + X_n) - E(X_1)\right| \geq \epsilon\right) = 0.$$

Beweis: Sei $S^{(n)} = \frac{X_1 + \dots + X_n}{n}$. Dann ist $E(S^{(n)}) = E(X_1)$, und

$$\begin{aligned} \text{Var}(S^{(n)}) &= \frac{1}{n^2} \text{Var}(X_1 + \dots + X_n) \\ &= \frac{1}{n^2} \cdot n \cdot M \\ &= \frac{M}{n}, \end{aligned}$$

wobei wir im vorletzten Schritt die Unabhängigkeit von $(X_i)_i$ verwendet haben. Die Behauptung folgt nun aus der Tschebyscheff-Ungleichung. $\square$

***Beispiel 9.1.56** (n-maliges Würfeln).

In Beispiel 9.1.44.3 hatten wir schon den Erwartungswert $E(X_i) = 3.5$ und in Beispiel 9.1.48 die Varianz für die Augenzahl beim einfachen Wurf des Laplace-Würfels berechnet. Wir betrachten nun zum n -fachen Wurf die gemittelte Summe $S^{(n)} = \frac{1}{n}(X_1 + \dots + X_n)$ der Augenzahlen. Nach dem schwachen Gesetz der großen Zahlen (Satz 9.1.55) ist zu einer vorgegebenen Schranke $\epsilon > 0$ bei häufigem Würfeln die Wahrscheinlichkeit, dass die beobachtete mittlere Augenzahl um mehr als ϵ von ihrem Erwartungswert $E(S^{(n)}) = 3.5$ abweicht klein, vorausgesetzt n ist hinreichend groß. Doch wie oft muss man z.B. würfeln, damit für $\epsilon = 0.1$ die Wahrscheinlichkeit einer Abweichung kleiner ist als 0.01? Solche Fragen werden wir noch in Kapitel 10.1.3 genauer betrachten. Hier geben wir mit einer sehr groben Abschätzung zufrieden, die auf der Tschebyscheff-Ungleichung (Satz 9.1.54) beruht, und wollen damit nur (9.36) an einem Beispiel illustrieren. Wir erhalten mit $M = \frac{35}{12}$ und $\epsilon = 0.1$:

$$P(|S^{(n)} - 3.5| \geq 0.1) \leq \frac{35}{12 \cdot 0.1 \cdot n}. \quad (9.37)$$

Die rechte Seite der Abschätzung (9.37) ist kleiner oder gleich 0.01, falls $n \geq 4200$. D.h. wenn man 4200 mal oder noch häufiger würfelt, dann weicht die mittlere Augenzahl mit einer Wahrscheinlichkeit von höchstens 1% um 0.1 oder mehr vom ihrem Erwartungswert ab.

***Bemerkung 9.1.57** (zum schwachen Gesetz der großen Zahlen).

Das schwache Gesetz der großen Zahlen sagt, dass in der Situation in Satz 9.1.55 für „große“ n der gemittelte Wert $S^{(n)} = \frac{1}{n}(X_1 + \dots + X_n)$ mit „großer“ Wahrscheinlichkeit (also einer

solchen nahe bei 1) vom Erwartungswert $E(S^{(n)}) = E(X_i)$ „nicht stark“ abweicht. Wenn man den Erwartungswert der Augenzahl bei einem Würfel statistisch durch viele Würfe ermitteln will, führt man aber z.B. *eine* recht lange Versuchsreihe von Würfeln durch, die einer Folge $X_1, X_2, \dots$ entspricht und betrachtet entsprechend die Folge der gemittelten Werte $S^{(1)}, S^{(2)}, \dots$. Das schwache Gesetz der großen Zahlen sagt, dass für ein vorgegebenes ϵ für hinreichend große n die Wahrscheinlichkeit für eine Abweichung $|S^{(n)} - E(X_1)| > \epsilon$ „klein“ ist, schließt aber nicht aus, dass für eine betrachtete Folge von Würfeln diese Abweichung „immer wieder mal“ auftritt. Aber das **starke Gesetz der großen Zahlen**, das wir hier nicht als mathematischen Satz formulieren, sagt, dass für *fast alle* Folgen (von Würfeln) die Folge der Werte von $S^{(n)}$ tatsächlich gegen $E(X_1)$ konvergiert. Das bedeutet, die Wahrscheinlichkeit für diese Konvergenz ist gleich 1.

9.2 Unendliche Wahrscheinlichkeitsräume

9.2.1 Diskrete Wahrscheinlichkeitsräume

Definition 9.2.1 (Diskreter Wahrscheinlichkeitsraum).

Seien Ω eine höchstens abzählbare Menge und $P : \mathcal{P}(\Omega) \rightarrow [0, 1]$ eine Funktion. Dann heißt (Ω, P) ein **diskreter Wahrscheinlichkeitsraum**, wenn folgendes gilt:

1.

$$P(\Omega) = 1. \quad (9.38)$$

2. Für jede Folge $A_1, A_2, \dots$ paarweiser disjunkter Teilmengen von Ω ist

$$P\left(\bigcup_{i=1}^{\infty} A_i\right) = \sum_{i=1}^{\infty} P(A_i). \quad (9.39)$$

Bemerkung 9.2.2. Eigenschaft (9.39) heißt **σ -Additivität**. Formal ist bei abzählbaren Wahrscheinlichkeitsräumen vieles analog zur Theorie der endlichen Wahrscheinlichkeitsräume (s. Kapitel 9.1). Nun ist aber bei der Summation (z.B. zur Berechnung des Erwartungswertes einer reellen Zufallsvariablen) die Summierbarkeit (absolute Konvergenz) i.a. nicht gewährleistet. Es gibt also reelle Wahrscheinlichkeitsverteilungen ohne endlichen Erwartungswert (s.u. Beispiel 9.2.3.2).

Beispiel 9.2.3 (für unendliche diskrete Wahrscheinlichkeitsräume).

1. (**Poisson-Verteilung**)

Eine bestimmte Masse einer radioaktiven Substanz zerfällt. Die Anzahl der Zerfälle $X_{[0,T]}$ im Zeitintervall $[0, T]$ ist eine Zufallsvariable. Dabei nehmen wir an, dass die Gesamtzahl der radioaktiven Teilchen sich im betrachteten Zeitraum nicht wesentlich ändert. Als mathematisches Modell nehmen wir die Verteilung

$$P_{\lambda}(X_{[0,T]} = k) = e^{-\lambda T} \frac{(\lambda T)^k}{k!} \quad \text{für } k \in \{0, 1, 2, \dots\}, \quad (9.40)$$

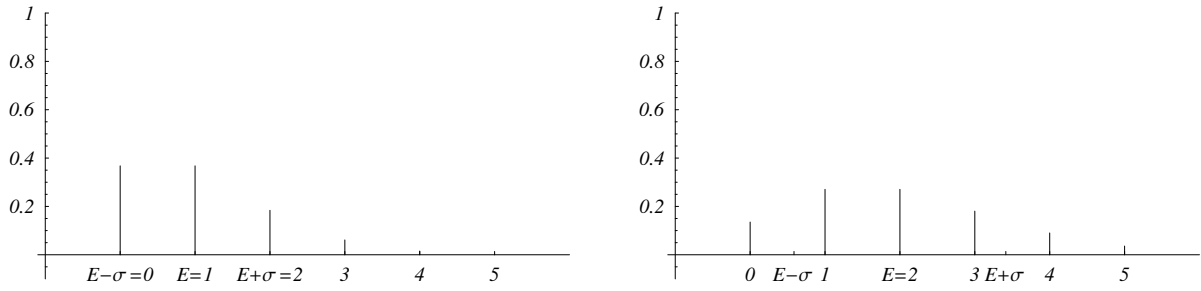


Abbildung 9.7: Stabdiagramme von Poisson-Verteilungen mit den Parametern $\lambda = 1$ und $T = 1$ (links), bzw. $T = 2$ (rechts)

mit einem Parameter $\lambda > 0$, die in Abbildung 9.7 illustriert ist. Es gilt für den Erwartungswert, das zweite Moment und die Varianz der Verteilung:

$$\begin{aligned}
 E(X_{[0,T]}) &= \sum_{k=0}^{\infty} k \cdot P_{\lambda}(X = k) \\
 &= \sum_{k=0}^{\infty} k e^{-\lambda T} \frac{(\lambda T)^k}{k!} \\
 &= \lambda T \cdot e^{-\lambda T} \sum_{k=1}^{\infty} \frac{(\lambda T)^{k-1}}{(k-1)!} \\
 &= \lambda T \cdot e^{-\lambda T} \sum_{l=0}^{\infty} \frac{(\lambda T)^l}{l!} \\
 &= \lambda T \cdot e^{-\lambda T} \cdot e^{\lambda T} \\
 &= \lambda T, \\
 E((X_{[0,T]})^2) &= \sum_{k=0}^{\infty} k^2 \cdot P_{\lambda}(X = k) \\
 &= \lambda T \cdot e^{-\lambda T} \sum_{k=1}^{\infty} k \frac{(\lambda T)^{k-1}}{(k-1)!} \\
 &= \lambda T \cdot e^{-\lambda T} \left[\sum_{k=1}^{\infty} (k-1) \frac{(\lambda T)^{k-1}}{(k-1)!} + \sum_{k=1}^{\infty} \frac{(\lambda T)^{k-1}}{(k-1)!} \right] \\
 &= \lambda T \cdot e^{-\lambda T} [\lambda T \cdot e^{\lambda T} + e^{\lambda T}] \\
 &= (\lambda T)^2 + \lambda T,
 \end{aligned}$$

$$\begin{aligned}\text{Var}(X_{[0,T]}) &= E((X_{[0,T]})^2) - (E(X_{[0,T]}))^2 \\ &= \lambda T.\end{aligned}$$

Des Weiteren gilt

$$\frac{dE(X_{[0,T]})}{dT} = \lambda,$$

d.h. λ ist die Zerfallsrate $\frac{\text{mittlere Anzahl der Zerfälle}}{\text{Zeit}}$.

2. **(Beispiel für eine Verteilung ohne endlichen Erwartungswert)** Wir betrachten die Zufallsvariable X mit der Verteilung

$$P(X = k) = \frac{6}{\pi^2} \cdot \frac{1}{k!} \quad \text{für } k \in \{1, 2, 3, \dots\}.$$

Es gilt

$$\begin{aligned}\sum_{k=1}^{\infty} P(X = k) &= \frac{6}{\pi^2} \sum_{k=1}^{\infty} \frac{1}{k^2} \\ &= 1.\end{aligned}$$

also handelt es sich tatsächlich um eine Wahrscheinlichkeitsverteilung. Aber wegen

$$\begin{aligned}E(X) &= \sum_{k=1}^{\infty} P(X = k) \cdot k \\ &= \frac{6}{\pi^2} \cdot \underbrace{\sum_{k=1}^{\infty} \frac{1}{k}}_{\text{divergente Reihe}} \\ &= \infty\end{aligned}$$

ist ihr Erwartungswert unendlich.

9.2.2 Kontinuierliche Wahrscheinlichkeitsräume

Wir betrachten nun den Fall, dass Ω ein Intervall ist, also z.B. $\Omega = [0, 1]$, $\Omega = [0, \infty]$ oder $\Omega =]-\infty, \infty[$. Für ein Wahrscheinlichkeitsmaß auf einer solchen Menge sollten ebenfalls die Axiome (9.38) und (9.39) wie bei diskreten Wahrscheinlichkeitsräumen (s. Definition 9.2.1) gelten. Allerdings ist es i.a. *nicht* möglich, für jede Teilmenge A von Ω die Wahrscheinlichkeit „ $P(A)$ “ zu definieren. Für einen strengen mathematischen Zugang muß man daher erst definieren, welche Teilmengen von Ω *meßbar* sind. Darauf gehen wir hier aber nicht ein. In diesem Abschnitt werden Begriffe nur heuristisch eingeführt. Wir geben also keine exakten Definitionen. Als Teilmengen A betrachten wir der Einfachheit halber nur Intervalle. Des Weiteren beschränken wir uns auf folgenden Spezialfall von Wahrscheinlichkeitsmaßen.

Definition 9.2.4. (Wahrscheinlichkeitsmaße mit einer Dichtefunktion)

Sei $\Omega = [a, b]$ ein Intervall mit $a < b$.

1. Eine **Wahrscheinlichkeitsdichte** auf Ω ist eine integrierbare Funktion $f : \Omega \rightarrow \mathbb{R}$ mit

(a)

$$f \geq 0,$$

d.h. $f(\omega) \geq 0$ für alle $\omega \in \Omega$.

(b)

$$\int_a^b f(\omega) d\omega = 1.$$

Die Wahrscheinlichkeitsdichte f ist also eine nicht-negative, **normierte** Funktion. Die Definition im Falle von (halb-) offenen Intervallen Ω sind analog.

2. Das zur Dichte f gehörende **Wahrscheinlichkeitsmaß** P ist auf Intervallen durch

$$P([a_0, b_0]) = \int_{a_0}^{b_0} f(\omega) d\omega \quad (9.41)$$

definiert, wie in Abbildung 9.8 illustriert.

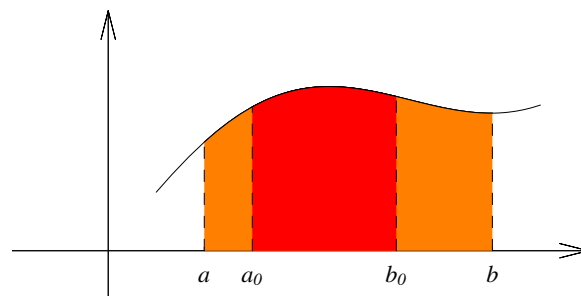


Abbildung 9.8: Wahrscheinlichkeitsdichte: Die Fläche über dem Intervall $[a_0, b_0]$ ist gleich der Wahrscheinlichkeit dieses Intervalls

3. Die Stammfunktion F von f , definiert durch

$$F(x) = \int_a^x f(\omega) d\omega,$$

heißt **Verteilungsfunktion** von P .

4. Eine **reelle Zufallsvariable** ist eine Funktion

$$X : \Omega \rightarrow \mathbb{R}.$$

Ihr **Erwartungswert** ist

$$E(X) := \int_a^b X(\omega) f(\omega) d\omega, \quad (9.42)$$

falls das Integral in (9.42) existiert, und ihre **Varianz** ist

$$\text{Var}(X) := \int_a^b (X(\omega) - E(X))^2 f(\omega) d\omega, \quad (9.43)$$

sofern die Integrale in (9.42) und (9.43) existieren.

Bemerkung 9.2.5. (Erwartungswert und Varianz einer Wahrscheinlichkeitsverteilung auf $\mathbb{R}$)

(Vgl. Bemerkung 9.1.41) Üblicherweise ist das durch P bestimmte Maß auf $\Omega = [a, b]$ schon das Bildmaß einer Funktion X mit Werten in $[a, b]$, wobei der Definitionsbereich von X nicht näher bekannt sein muß. Wir bezeichnen daher mit

$$\mu = \int_a^b x \cdot f(x) dx \quad (9.44)$$

den **Erwartungswert der Verteilung** und mit

$$\sigma^2 = \int_a^b (x - \mu)^2 f(x) dx \quad (9.45)$$

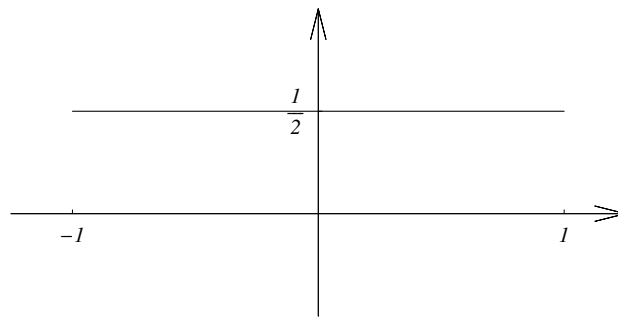
ihre **Varianz**, sofern diese Integrale existieren.

***Bemerkung 9.2.6.** Formal kann man den Bezug zwischen (9.44), bzw. (9.45) zur Definition des Erwartungswertes, bzw. der Varianz einer Zufallsvariablen (s. (9.42) bzw. (9.43)) herstellen, indem man den Erwartungswert (die Varianz) einer reellen Verteilung als den Erwartungswert (die Varianz) der durch $X(x) = x$ definierten Zufallsvariablen betrachtet.

Beispiel 9.2.7. (Gleichverteilung auf einem beschränkten Intervall)

Die **Gleichverteilung** auf $[a, b]$ ist durch die Dichtefunktion

$$f : [a, b] \rightarrow \mathbb{R}, \\ x \mapsto \frac{1}{b - a},$$

Abbildung 9.9: Gleichverteilung auf dem Intervall $[-1, 1]$

gegeben (s. Abbildung 9.9.)

Es gelten

$$f(x) = \frac{1}{b-a} > 0$$

und

$$\int_a^b f(x) dx = 1,$$

d.h. f ist also tatsächlich eine Wahrscheinlichkeitsdichte. Zur Vereinfachung der Notation betrachten wir eine Zufallsvariable X , deren Verteilung die Dichte f hat. (Dann können wir nämlich für die im Folgenden betrachteten Erwartungswerte $E(X)$, $E(X^2)$ etc. schreiben.) Der Erwartungswert der Verteilung ist

$$\begin{aligned} E(X) &= \int_a^b \frac{1}{b-a} \cdot x dx \\ &= \frac{1}{b-a} \cdot \frac{1}{2}(b^2 - a^2) \\ &= \frac{b+a}{2}, \end{aligned}$$

also gleich dem Mittelpunkt des Intervalls $[a, b]$. Zur Berechnung der Varianz benutzen wir

$$\begin{aligned} \text{Var}(X) &= E((X - E(X))^2) \\ &= E(X^2) - (E(X))^2. \end{aligned}$$

Wir müssen also noch das **zweite Moment** $E(X^2)$ von X berechnen.

$$\begin{aligned} E(X^2) &= \int_a^b \frac{1}{b-a} x^2 dx \\ &= \frac{1}{b-a} \cdot \frac{1}{3} (b^3 - a^3) \\ &= \frac{1}{3} (b^2 + ab + a^2). \end{aligned}$$

Damit erhalten wir

$$\begin{aligned} \text{Var}(X) &= \frac{1}{3} (b^2 + ab + a^2) - \frac{1}{4} (b^2 + 2ab + a^2) \\ &= \frac{1}{12} (b-a)^2. \end{aligned}$$

Die Varianz hängt also nur von der Intervalllänge ab. Physikalisch kann man den Erwartungswert von X als *Schwerpunkt* bei homogener Massenverteilung interpretieren, und die Varianz ist proportional zum *Trägheitsmoment*, also proportional zum *mittleren quadratischen Abstand* zum Schwerpunkt.

Beispiel 9.2.8 (Exponentialverteilungen auf $[0, \infty)$).

Die **Exponentialverteilung** mit Parameter $\lambda > 0$ ist gegeben durch die Dichte

$$\begin{aligned} f_\lambda : [0, \infty) &\rightarrow \mathbb{R}, \\ r &\mapsto \lambda e^{-\lambda t}. \end{aligned}$$

Sie tritt z.B. beim durch den Poisson-Prozeß modellierten radioaktiven Zerfall auf (s. Beispiel 9.2.3.1.) Die **Wartezeit** bis zum ersten Zerfall (nach einem festgelegten Zeitpunkt, den wir hier als 0 wählen) ist eine Zufallsvariable, deren Verteilung die Dichte f_λ hat. Die Wahrscheinlichkeit dafür, dass nach der Zeitdauer T noch kein Zerfall eingetreten ist, ist gleich

$$\begin{aligned} P_\lambda((T, \infty)) &= \int_T^\infty \lambda e^{-\lambda t} dt \\ &= [-e^{-\lambda t}]_T^\infty \\ &= e^{-\lambda T}. \end{aligned}$$

Dies ist gleich der Wahrscheinlichkeit $P_\lambda(X_{[0,T]} = 0)$ der Poisson-Verteilung (9.40).

Wir weisen nun nach, dass f_λ eine Wahrscheinlichkeitsdichte ist und berechnen den Erwartungswert, das zweite Moment und die Varianz der Verteilung aus: Die Funktion f_λ nimmt offensichtlich nur positive Werte an und ist wegen

$$\begin{aligned} \int_0^\infty \lambda \cdot e^{-\lambda x} dx &= -e^{-\lambda x} \Big|_0^\infty \\ &= 1 \end{aligned}$$

normiert, also eine Wahrscheinlichkeitsdichte. Der Erwartungswert ist

$$\begin{aligned}
 \mu &= \int_0^{\infty} x \cdot \lambda \cdot e^{-\lambda x} dx \quad (\text{partielle Integration}) \\
 &= \underbrace{-xe^{-\lambda x}}_{=0} \Big|_0^{\infty} + \int_0^{\infty} e^{-\lambda x} dx \\
 &= -\frac{1}{\lambda} e^{-\lambda x} \Big|_0^{\infty} \\
 &= \frac{1}{\lambda}.
 \end{aligned}$$

Das zweite Moment der Verteilung ist

$$\begin{aligned}
 \int_0^{\infty} x^2 \cdot \lambda \cdot e^{-\lambda x} dx &= \underbrace{-x^2 e^{-\lambda x}}_{=0} \Big|_0^{\infty} + 2 \int_0^{\infty} x e^{-\lambda x} dx \quad (\text{durch partielle Integration}) \\
 &= 2 \cdot \frac{1}{\lambda} \cdot \frac{1}{\lambda} \\
 &= \frac{2}{\lambda^2}.
 \end{aligned}$$

Also ist die Varianz gleich

$$\begin{aligned}
 \sigma^2 &= \frac{2}{\lambda^2} - \left(\frac{1}{\lambda}\right)^2 \\
 &= \frac{1}{\lambda^2}.
 \end{aligned}$$

Beispiel 9.2.9 (Normalverteilungen).

Die **Normalverteilung** $\mathcal{N}(\mu, \sigma^2)$ mit Erwartungswert μ und Varianz σ^2 hat die Dichte

$$f_{\mu, \sigma^2}(x) = \frac{1}{\sigma \sqrt{2\pi}} e^{\left(\frac{-(x-\mu)^2}{2\sigma^2}\right)}. \quad (9.46)$$

Die Normalverteilung $\mathcal{N}(0, 1)$ mit Erwartungswert 0 und Varianz 1 heißt **Standard-Normalverteilung**.

Durch die Normalverteilung werden viele gestreute Größen, wie z.B. Körperlängen von Personen in einer Bevölkerung beschrieben, allerdings nur in einem hinreichend kleinen Intervall um die Durchschnittsgröße herum, denn natürlich gibt es keinen Menschen mit negativer Größe oder von 3m Länge. Solche Verteilungen haben mit den Normalverteilungen die typische Glockenform gemeinsam. Mathematisch wird der Zustand zwischen der Normalverteilung und mehrfach wiederholten Experimenten (z.B. mehrfacher Münzwurf) durch den **zentralen Grenzwertsatz**

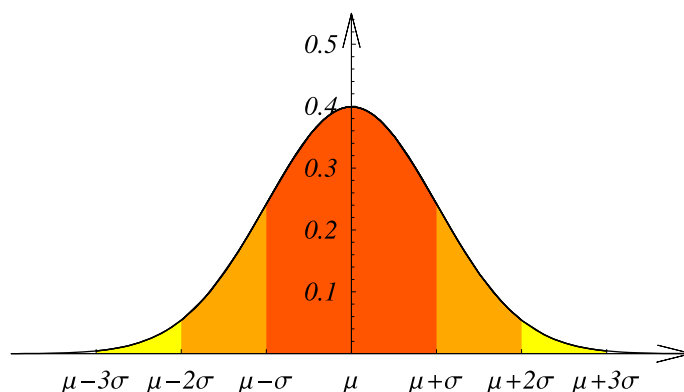


Abbildung 9.10: Die Standard-Normalverteilung mit ihrem σ -, 2σ - und 3σ -Intervall

(Satz 9.2.10) hergestellt.

Wir überprüfen die Normiertheit und berechnen den Erwartungswert und die Varianz. Zunächst sehen wir (z.B. mit Hilfe des Majorantenkriteriums), dass das uneigentliche Integral

$$I := \int_{-\infty}^{\infty} e^{-x^2} dx \quad (9.47)$$

existiert. Zu der Funktion e^{-x^2} gibt es keine *elementare* Stammfunktion, wie wir schon in Bemerkung 4.10.31 erwähnt hatten. Dennoch können wir den Wert von I exakt berechnen, und zwar mit Hilfe von Integration in 2d und Polarkoordinaten (vgl. Abschnitt 7.8.2, Beispiel 7.8.3). Es gilt nämlich

$$\begin{aligned} I^2 &= \int_{-\infty}^{\infty} e^{-x^2} dx \cdot \int_{-\infty}^{\infty} e^{-y^2} dy \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-x^2-y^2} dx dy. \end{aligned}$$

Wir substituieren: $x = r \cos \varphi$, $y = r \sin \varphi$, $dx dy = r d\varphi dr$, und erhalten damit

$$\begin{aligned}
 I^2 &= \int_0^\infty \int_0^{2\pi} e^{-r^2} r d\varphi dr \\
 &= 2\pi \int_0^\infty r e^{-r^2} dr \\
 &= \pi \int_0^\infty 2r e^{-r^2} dr \\
 &= \pi [-e^{-r^2}]_0^\infty \\
 &= \pi.
 \end{aligned}$$

Also ist $I = \sqrt{\pi}$. In der folgenden Rechnung verwenden wir die Substitution

$$\begin{aligned}
 y &= \frac{x - \mu}{\sqrt{2}\sigma}, \\
 \Leftrightarrow x &= \sqrt{2}\sigma y + \mu, \\
 dx &= \sqrt{2}\sigma dy,
 \end{aligned}$$

und erhalten die Normiertheit der Dichtefunktion:

$$\begin{aligned}
 \int_{-\infty}^\infty \frac{1}{\sigma\sqrt{2\pi}} e^{\left(\frac{-(x-\mu)^2}{2\sigma^2}\right)} &= \int_{-\infty}^\infty \frac{1}{\sigma\sqrt{2\pi}} \cdot \sqrt{2}\sigma e^{-y^2} dy \\
 &= \frac{1}{\sqrt{\pi}} \int_{-\infty}^\infty e^{-y^2} dy \\
 &= 1.
 \end{aligned}$$

Zur Berechnung des Erwartungswertes einer $\mathcal{N}(\mu, \sigma^2)$ -verteilten Zufallsvariablen X_{μ, σ^2} (die Verteilung dieser Zufallsvariablen hat also die Dichte f_{μ, σ^2}) verwenden wir die Symmetrie von f_{μ, σ^2} , d.h. die Identität

$$f_{\mu, \sigma^2}(\mu + y) = f_{\mu, \sigma^2}(\mu - y) \quad \forall y \in \mathbb{R},$$

sowie die Substitution $x = y + \mu$ und $x = -y + \mu$ im ersten und zweiten Integral in (9.48),

respektive.

$$\begin{aligned}
 E(X_{\mu,\sigma^2}) &= \int_{-\infty}^{\infty} x \cdot f_{\mu,\sigma^2}(x) dx \\
 &= \int_{-\infty}^{\mu} x f_{\mu,\sigma^2}(x) dx + \int_{\mu}^{\infty} x \cdot f_{\mu,\sigma^2}(x) dx \\
 &= \int_{-\infty}^0 (y\mu) f_{\mu,\sigma^2}(y) dy + \int_{-\infty}^0 (-y + \mu) f_{\mu,\sigma^2}(y) dy \quad (9.48) \\
 &= \mu \cdot 2 \int_{-\infty}^{\infty} f_{0,\sigma^2}(y) dy \\
 &= \mu \int_{-\infty}^{\infty} f_{0,\sigma^2}(y) dy \\
 &= \mu.
 \end{aligned}$$

Wir haben schon mehrfach bemerkt, dass die Varianz invariant bezüglich einer „Verschiebung“ der Dichte ist, d.h. für jedes $v \in \mathbb{R}$ haben zwei Verteilungen mit Dichten $f(\cdot)$ und $f(\cdot - v)$ die gleiche Varianz. Wir berechnen nun die Varianz der zentrierten Verteilungen unter Verwendung der Substitution $y = \sqrt{2}\sigma x$.

$$\begin{aligned}
 \text{Var}(X_{0,\sigma^2}) &= \int_{-\infty}^{\infty} x^2 \frac{1}{\sigma\sqrt{2\pi}} e^{\left(\frac{-x^2}{2\sigma^2}\right)} dx \\
 &= \int_{-\infty}^{\infty} \frac{2\sigma^2 y^2}{\sigma\sqrt{2\pi}} e^{-y^2} \cdot \sqrt{2}\sigma dy \\
 &= \frac{2\sigma^2}{\sqrt{\pi}} \int_{-\infty}^{\infty} y^2 e^{-y^2} dy \\
 &= \frac{2\sigma^2}{\sqrt{\pi}} \cdot \frac{-1}{2} \cdot \int_{-\infty}^{\infty} y \cdot (-2y \cdot e^{-y^2}) dy \\
 &= \frac{-\sigma}{\sqrt{\pi}} \left[\underbrace{[y \cdot e^{-y^2}]_{-\infty}^{\infty}}_{=0} - \int_{-\infty}^{\infty} e^{-y^2} dy \right] \\
 &= \sigma^2.
 \end{aligned}$$

Dabei haben wir im vorletzten Schritt partiell integriert.

Der zentrale Grenzwertsatz, den wir hier in einer speziellen Version formulieren, erklärt die herausragende Bedeutung von Normalverteilungen für die Wahrscheinlichkeitstheorie und Statistik.

Satz 9.2.10 (Zentraler Grenzwertsatz).

Sei $X_1, X_2, \dots$ eine Folge von auf demselben Wahrscheinlichkeitsraum (Ω, P) definierten, paarweise unabhängigen, identisch verteilten reellen Zufallsvariablen mit

$$E(X_i) = \mu, \quad \text{Var}(X_i) = \sigma^2 > 0.$$

Sei

$$X^{(n)} = X_1 + \dots + X_n,$$

und sei

$$Z^{(n)} = \frac{X^{(n)} - n\mu}{\sigma\sqrt{n}}.$$

(Wir erhalten $Z^{(n)}$ also aus $X^{(n)}$ durch Zentrierung und Standardisierung. Somit hat $Z^{(n)}$ den Erwartungswert 0 und die Varianz 1.) Dann gilt für jedes Intervall $[a_0, b_0] \subset \mathbb{R}$:

$$\lim_{n \rightarrow \infty} P(Z^{(n)} \in [a_0, b_0]) = \int_{a_0}^{b_0} f_{0,1}(x) dx.$$

wobei $f_{0,1}$ die Dichte der Standard-Normalverteilung ist. Äquivalent dazu können wir schreiben:

$$\lim_{n \rightarrow \infty} P\left(\frac{X^{(n)} - n\mu}{\sigma\sqrt{n}} \in [a_0, b_0]\right) = \int_{a_0}^{b_0} f_{0,1}(x) dx.$$

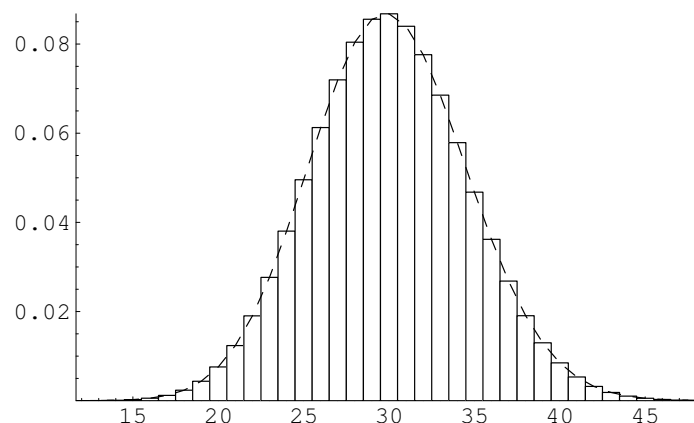


Abbildung 9.11: Histogramm der Binomialverteilung für $n = 100$ und $p = 0.3$, verglichen mit der $\mathcal{N}(np, np(1 - p))$ Verteilung.

Beispiel 9.2.11 (Binomialverteilung für große n).

Die Binomialverteilung mit gegebenem Erfolgsparameter p wird für große n ungefähr gleich einer $\mathcal{N}(np, np(1-p))$ Normalverteilung:

$$P(k) = \binom{n}{k} p^k (1-p)^{n-k} \approx \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(k-\mu)^2}{2\sigma^2}} \text{ mit } \mu = np \text{ und } \sigma^2 = np(1-p).$$

Dieser Sachverhalt, der für $p = 0.3$ und $n = 100$ in Abbildung 9.11 illustriert ist, folgt direkt aus dem zentralen Grenzwertsatz, denn die binomialverteilte Zufallsvariable K kann als Summe vieler unabhängiger Zufallsvariablen X_i aufgefasst werden, die jeweils nur die Werte 0 oder 1 (jeweils mit Wahrscheinlichkeit $(1-p)$ bzw. p) annehmen, und die den Erwartungswert p und die Varianz $p(1-p)$ haben.

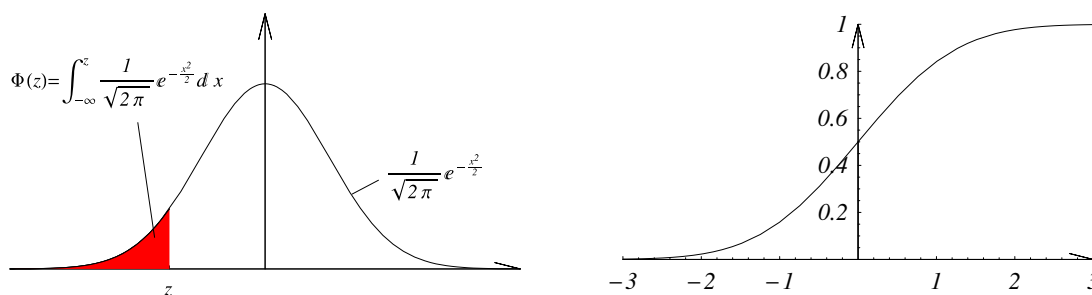


Abbildung 9.12: Die Standard-Normalverteilung (links) und ihre Verteilungsfunktion (rechts)

Definition 9.2.12. (Verteilungsfunktion der Standard-Normalverteilung)

Die **Verteilungsfunktion** (s. Definition 9.2.4.3) **der Standard-Normalverteilung** ist

$$\begin{aligned} \Phi : \mathbb{R} &\rightarrow \mathbb{R}, \\ \Phi(z) &= \int_{-\infty}^z f_{0,1}(x) dx. \end{aligned}$$

Die Graphen der Dichte $f_{0,1}$ und von Φ ist in Abbildung 9.12 zu sehen.

Bemerkung 9.2.13. (zur Verteilungsfunktion der Standard-Normalverteilung)

Bekanntlich gibt es keine Darstellung von Φ durch *elementare* Funktionen (s. Bemerkung 4.10.31.) Werte von Φ lassen sich aber beliebig genau numerisch berechnen und für diskrete Werte von z liegen die Funktionswerte tabellarisch vor, wodurch man schnell Integrale

$$\int_a^b f_{0,1}(x) dx = \Phi(b) - \Phi(a)$$

durch das Auswerten von Φ an den Integrationsgrenzen beliebig genau auswertet. Wegen

$$\Phi(-z) = 1 - \Phi(z)$$

enthalten solche Tabellen z.B. nur die Werte für nicht-negative z .

Mit der folgenden nützlichen Identität kann man die Wahrscheinlichkeit von Intervallen $[-z, z]$ (mit $z > 0$) ausrechnen, die symmetrisch bzgl. des Erwartungswertes 0 der Normalverteilung sind.

$$\begin{aligned} \int_{-z}^z f_{0,1}(x) dx &= \Phi(z) - \Phi(-z) \\ &= \Phi(z) - (1 - \Phi(z)) \\ &= 2\Phi(z) - 1. \end{aligned}$$

Einige spezielle Werte von Φ und oder die der entsprechenden Integrale sollten allen Anwendern statistischer Methoden bekannt sein:

$$\begin{aligned} \Phi(0) &= 0.5, \\ \Phi(1) &\approx 0.8413 \quad \Rightarrow \quad \int_{-1}^1 f_{0,1}(y) dy \approx 0.6826, \\ \Phi(1.96) &\approx 0.975 \quad \Rightarrow \quad \int_{-1.96}^{1.96} f_{0,1}(y) dy \approx 0.95, \\ \Phi(2) &\approx 0.9772 \quad \Rightarrow \quad \int_{-2}^2 f_{0,1}(y) dy \approx 0.9544, \\ \Phi(3) &\approx 0.9986 \quad \Rightarrow \quad \int_{-3}^3 f_{0,1}(y) dy \approx 0.9972. \end{aligned} \tag{9.49}$$

Aus der zweiten Zeile folgt z.B., dass bei irgendeiner Normalverteilung dem Intervall $[\mu - \sigma, \mu + \sigma]$ mit Radius σ (Streuung) um den Erwartungswert μ herum eine Wahrscheinlichkeit von etwa 68% zugeordnet wird (vgl. dazu Abbildung 9.12.) Bei einem Experiment mit vielen voneinander unabhängigen $\mathcal{N}(\mu, \sigma^2)$ -verteilten Messungen liegen ungefähr 68% der Meßwerte in diesem Intervall.

Definition 9.2.14 (α -Quantile der $\mathcal{N}(\mu, \sigma^2)$ -Verteilung).

Sei $\alpha \in]0, 1[$. Das α -**Quantil** der Standard-Normalverteilung ist die Zahl $z \in \mathbb{R}$ mit

$$\alpha = \int_{-\infty}^z f_{0,1}(x) dx = \Phi(z),$$

also

$$z = \Phi^{-1}(\alpha).$$

Bemerkung 9.2.15. (Quantile für allgemeine Verteilungen, Median)

Wir erwähnen noch, dass man α -Quantile allgemein für (diskrete oder kontinuierliche) reelle Verteilungen definieren kann, was wir hier aber wegen der dafür nötigen Fallunterscheidungen nicht tun. Das $\frac{1}{2}$ -Quantil heißt **Median** der Verteilung. Im Falle einer kontinuierlichen Verteilung auf einem Intervall $[a, b]$ mit überall positiver Dichte f ist der Median m die durch die Bedingung $P([a, m]) = \frac{1}{2}$ eindeutig festgelegte Zahl. Der Median ist im Allgemeinen vom Erwartungswert verschieden.

Kapitel 10

Statistik

In diesem Kapitel können wir nur einige Ideen der für Anwendungen so wichtigen Statistik vorstellen und hoffen, dass unsere Vorgehensweise, erst die Wahrscheinlichkeitstheorie als Grundlage für ein tieferes Verständnis der Statistik relativ ausführlich behandelt zu haben, dem Leser spätestens im Nachhinein gerechtfertigt erscheint. Den sicheren Gebrauch statistischer Methoden lernt man am besten durch Anwendung. Hierfür gibt es im dritten Semester eine spezielle Veranstaltung.

Als wichtigste Quelle zur Vorlesungsvorbereitung zu diesem Kapitel diene [Kre02]. Eine elementare Einführung in die Statistik ist [Bos00]. Als Referenz für statistische Datenanalyse mit vielen anwendungsorientierten Beispielen möchten wir noch [Sac02] und [Sta02] nennen sowie das auf biologische Anwendungen ausgerichtete Standardwerk [SR94]. Unterhaltsam und informativ sind die eher populärwissenschaftlichen Bücher [BBDH01] und [Krä00], die insbesondere den falschen Gebrauch von Statistik illustrieren.

10.1 Parameterschätzung

In naturwissenschaftlichen Experimenten geht es insbesondere darum, von den gemachten Beobachtungen auf charakteristische Größen eines Systems zu schließen. In manchen Fällen sind solche Größen „direkt“ messbar, z.B. die Länge eines bestimmten Metallstabs unter bestimmten Bedingungen (z.B. Temperatur). Mehrmaliges Messen sollte idealerweise stets zum gleichen Ergebnis führen. Unterliegt jedoch die Messung zufälligen Schwankungen aufgrund nicht auszuschließender Ungenauigkeiten der Messapparatur oder sind die beobachteten Größen selber zufällig verteilt, wie z.B. die Anzahl der radioaktiven Zerfälle pro Sekunde einer bestimmten Testsubstanz, so können wir die Messungen/Beobachtungen als Ausgang (**Realisierung** oder **Stichprobe**) eines Zufallsexperiments auffassen.

Zur Interpretation der Beobachtungen gehen wir von möglichen Modellen für das beobachtete System aus, d.h. wir betrachten die Menge aller möglichen Ausgänge eines Experiments und auf dieser Menge verschiedene Wahrscheinlichkeitsmaße. Diese sind üblicherweise durch einen Parameter gekennzeichnet. Dieser kann z.B. durch ein n -Tupel von reellen Zahlen gegeben sein. Bei Kenntnis des Wertes dieses Parameters wüßte man also die (diesem Parameterwert zuge-

ordnete) Verteilung und hätte somit das Zufallsexperiment vollständig durch einen Wahrscheinlichkeitsraum beschrieben. Von einer solchen Kenntnis sind wir in Kapitel 9 stets ausgegangen und konnten so allen Ereignissen eine Wahrscheinlichkeit zuordnen. Nun ist aber der Wert des Parameters und somit das Wahrscheinlichkeitsmaß unbekannt.

Die Aufgabe besteht darin, aufgrund der Kenntnis von Realisierungen den Parameter zu schätzen, also allgemein einen **Schätzer** anzugeben, also eine Vorschrift, die jeder möglichen Stichprobe (Ausgang des Zufallsexperiments) einen Parameterwert zuordnet. Die Wahl eines solchen Schätzers ist keineswegs durch das Zufallsexperiment und den zu schätzenden Parameter eindeutig vorgegeben. Oft bieten sich verschiedene Schätzer an. Wir stellen hier exemplarisch einige solcher Schätzer zu uns aus Kapitel 9 bereits bekannten Zufallsexperimenten vor und beschreiben einige ihrer Eigenschaften und somit mögliche Auswahlkriterien.

10.1.1 Schätzprobleme und Schätzer

Beispiel 10.1.1 (Erfolgparameter bei einem Münzwurf).

Wir betrachten eine Münze mit unbekanntem Erfolgparameter p , der Wahrscheinlichkeit für das Ereignis „Kopf“. Dazu definieren wir für den i -ten Münzwurf die reelle Zufallsvariable X_i , die bei dem Ereignis „ i -ter Wurf ist Kopf“ den Wert 1 annimmt und sonst den Wert 0. Die X_i sind also voneinander unabhängig und identisch verteilt mit

$$\begin{aligned} P_p(X_i = 1) &= p, \\ P_p(X_i = 0) &= 1 - p. \end{aligned}$$

Durch die Indizierung P_p deuten wir an, dass das Wahrscheinlichkeitsmaß von dem Parameter p abhängt, dessen numerischer Wert uns nun nicht bekannt ist. Der Erwartungswert der Verteilung von X_i ist $E(X_i) = p$. Ein Experiment von n auf einanderfolgenden Münzwürfen entspricht der Zufallsvariable $X = (X_1, \dots, X_n)$. Mit ihnen können wir die Zufallsvariable

$$X^{(n)} := \frac{1}{n}(X_1 + \dots + X_n) \quad (10.1)$$

definieren, also die durchschnittliche Anzahl der Erfolge (Achtung: Durchschnitt bedeutet hier „Division durch die Anzahl n der Würfe“, also die Bildung des arithmetischen Mittels und ist *nicht* mit dem Erwartungswert zu verwechseln.)

Wir möchten nun den Erfolgparameter p **schätzen**. Es erscheint intuitiv sinnvoll, jeder Realisierung $x = (x_1, \dots, x_n) = (X_1(\omega), \dots, X_n(\omega))$ den folgenden Schätzwert von p zuzuordnen:

$$\hat{p}(x) := \frac{1}{n}(x_1 + \dots + x_n) \quad (10.2)$$

also das **arithmetische Mittel** der x_i oder, anders formuliert, die **relative Häufigkeit** der beobachteten Erfolge.

Achtung: Der **Schätzer** ist eine Funktion auf der Menge $\chi_1 \times \dots \times \chi_n$ der Realisierungen. Er ordnet jeder Realisierung x einen Schätzwert für den Parameter p zu. Manchmal wird auch kurz, aber strenggenommen nicht ganz korrekt, nur $\hat{p}$ anstatt $\hat{p}(x)$ geschrieben.

Der Schätzwert hängt von der jeweiligen Realisierung ab und diese ist zufällig. Diese Verknüpfung von Schätzer und der Zufallsvariable $X_1 \times \dots \times X_n$ ist gerade die in (10.1) definierte Zufallsvariable $X^{(n)}$.

Wir rechnen leicht nach, dass diese Zufallsvariable den gleichen Erwartungswert p hat, also gerade den Wert des zu schätzenden Parameters:

$$\begin{aligned} E_p(X^{(n)}) &= \frac{1}{n}(E_p(X_1) + \dots + E_p(X_n)) \\ &= \frac{1}{n}(p + \dots + p) \\ &= p. \end{aligned} \tag{10.3}$$

Nach diesem Beispiel geben wir die Definitionen der bereits illustrierten Begriffe.

Definition 10.1.2 (Schätzproblem).

Ein **Schätzproblem mit endlichem Stichprobenraum** ist durch folgendes gegeben.

1. Eine nicht-leere, endliche Menge χ , den **Stichprobenraum**,
2. eine Familie $\{P_\vartheta : \vartheta \in \Theta\}$ von Wahrscheinlichkeitsmaßen auf χ ,
3. einen **zu schätzenden Parameter** $g(\vartheta)$, wobei g eine Funktion auf Θ ist.

Definition 10.1.3 (Schätzer).

Sei Y der Wertebereich von g aus Definition 10.1.2. Dann ist jede Funktion

$$T : \chi \rightarrow Y$$

ein **Schätzer** von $g(\vartheta)$.

Beispiel 10.1.4. (Anwendung der Definitionen auf den n-fachen Münzwurf)

In Beispiel 10.1.1 ist der Stichprobenraum die Menge $\chi = \{0, 1\}^n$ aller binären n -Tupel. Die betrachteten Maße auf χ sind die Produktmaße P_p , die sich aus den jeweiligen Verteilungen auf $\{0, 1\}$ zum Parameter p ergeben (vgl. hierzu Beispiel 9.1.24). Der Parameter der Familie von Maßen ist $\vartheta = p$. Und da dieser selber geschätzt werden soll ist, ist $g(\vartheta) = \vartheta$.

10.1.2 Eigenschaften von Schätzern

Eine oftmals wünschenswerte Eigenschaft eines Schätzers haben wir in 10.1.1 bereits kennengelernt und in (10.3) für den dort betrachteten Schätzer nachgewiesen.

Definition 10.1.5. (Erwartungswert und Erwartungstreue eines Schätzers)

1. Zu einem gegebenen Schätzproblem (s. Definition 10.1.2) ist für jedes $\vartheta \in \Theta$ ein Wahrscheinlichkeitsraum (χ, P_ϑ) definiert und wir können auf diesem einen reellwertigen Schätzer T als reelle Zufallsvariable betrachten. Somit ist insbesondere zu jedem $\vartheta \in \Theta$ der **Erwartungswert des Schätzers bezüglich P_ϑ** definiert, und zwar durch

$$E_\vartheta(T) = \sum_{x \in \chi} T(x) P_\vartheta(x).$$

2. Ein Schätzer heißt **erwartungstreu**, wenn für jedes $\vartheta \in \Theta$ sein Erwartungswert bzgl. P_ϑ mit dem zu schätzenden Parameter $g(\vartheta)$ übereinstimmt, also

$$E_\vartheta(T) = g(\vartheta).$$

Beispiel 10.1.6 (Erwartungstreue Schätzung des Erwartungswertes). Wir verallgemeinern unsere Betrachtungen zur erwartungstreuen Schätzung des Erwartungswertes aus Beispiel 10.1.1. Sei also $X_1, X_2, \dots$ eine Folge von identisch verteilten Zufallsvariablen auf einem nicht genauer bekanntem Wahrscheinlichkeitsraum (Ω, P) , mit Werten in einer endlichen Menge $\chi_1 \subset \mathbb{R}$. Sei $E(X_i) = \mu$. Die einzelnen Zufallsvariablen können z.B. die Augenzahl beim Würfeln beschreiben. Dann ist μ gerade der Erwartungswert für die Augenzahl bei einem Wurf.

Bei n -fachem Wurf erhalten wir n -Tupel von Augenzahlen, also Werte $x = (x_1, \dots, x_n) \in \chi_1^n =: \chi$.

Wir definieren nun den Schätzer

$$\begin{aligned} T_0 : \chi &\rightarrow \mathbb{R}, \\ (x_1, \dots, x_n) &\mapsto \frac{1}{n}(x_1 + \dots + x_n). \end{aligned} \tag{10.4}$$

Also jeder Realisierung x wird das arithmetische Mittel als Schätzwert zugeordnet. Dieser Schätzer ist erwartungstreu, denn völlig analog zu (10.3) gilt

$$\begin{aligned} E(T_0) &= \frac{1}{n}(E(X_1) + \dots + E(X_n)) \\ &= \frac{1}{n}(\mu + \dots + \mu) \\ &= \mu. \end{aligned} \tag{10.5}$$

Wir bemerken, dass wir keine Voraussetzungen an die Unabhängigkeit der X_i gemacht haben. Des Weiteren gelten unsere Betrachtungen gleichfalls für Zufallsvariablen mit abzählbar diskreten oder kontinuierlichen Verteilungen, sofern deren Erwartungswert existiert. Der hier betrachtete Schätzer ist also z.B. auch für physikalische Messreihen geeignet, bei denen eine Messung durch eine kontinuierliche Wahrscheinlichkeitsverteilung modelliert wird.

Beispiel 10.1.7 (Erwartungstreue Schätzung der Varianz).

Wir untersuchen nun in der gleichen Situation wie in Beispiel 10.1.6 verschiedene Schätzer für die Varianz $\sigma^2 = \text{Var}(X_i)$ bei insgesamt n -fach durchgeführtem Experiment, das durch die Zufallsvariable $X_1 \times \dots \times X_n$ mit Werten in $\chi = \chi_1^n \subset \mathbb{R}^n$ beschrieben ist.

1. Wir nehmen zunächst an, der Erwartungswert $\mu = E(X_i)$ sei uns bekannt. Dann können wir den folgenden Schätzer definieren:

$$\begin{aligned} T_1 : \chi &\rightarrow \mathbb{R}, \\ (x_1, \dots, x_n) &\mapsto \frac{1}{n} \sum_{i=1}^n (x_i - \mu)^2. \end{aligned}$$

Dieser Schätzer ist in der Tat erwartungstreu, denn

$$\begin{aligned} E(T_1) &= \frac{1}{n} \sum_{i=1}^n E((X_i - \mu)^2) \\ &= \frac{1}{n} \sum_{i=1}^n \text{Var}(X_i) \\ &= \frac{1}{n} \sum_{i=1}^n \sigma^2 \\ &= \sigma^2. \end{aligned}$$

2. Im Allgemeinen ist uns der Erwartungswert μ aber nicht bekannt und wir müssen diesen auch schätzen. Dazu verwenden wir T_0 aus (10.4). Ein naheliegender Versuch für einen Schätzer der Varianz ist

$$\begin{aligned} T_2 : \chi &\rightarrow \mathbb{R}, \\ x = (x_1, \dots, x_n) &\mapsto \frac{1}{n} \sum_{i=1}^n (x_i - T_0(x))^2. \end{aligned}$$

Wir betrachten jetzt wieder $x_1, \dots, x_n$ als Werte der Zufallsvariablen $X_1, \dots, X_n$, respektive. Mit der Notation $\bar{X} = \frac{1}{n}(X_1 + \dots + X_n)$ können wir dann

$$T_2(X_1, \dots, X_n) = \frac{1}{n} \sum_{k=1}^n (X_k - \bar{X})^2$$

als Zufallsvariable auffassen und deren Erwartungswert, also den Erwartungswert des Schätzers T_2 ausrechnen. Dazu machen wir erst folgende Nebenrechnungen, in denen μ den unbekannten tatsächlichen Erwartungswert der X_i bezeichnet und σ^2 ihre tatsächliche

Varianz.

$$\begin{aligned}
 E((\bar{X} - \mu)^2) &= E\left(\left(\frac{1}{n} \sum_{i=1}^n X_i - \mu\right)^2\right) \\
 &= E\left(\left(\frac{1}{n} \sum_{i=1}^n (X_i - \mu)\right)^2\right) \\
 &= \frac{1}{n^2} E\left(\left(\sum_{i=1}^n (X_i - \mu)\right)^2\right) \\
 &= \frac{1}{n^2} E\left(\sum_{i=1}^n \sum_{j=1}^n (X_i - \mu)(X_j - \mu)\right) \\
 &= \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n \text{Cov}(X_i, X_j) \\
 &= \frac{1}{n} \sigma^2.
 \end{aligned}$$

Dabei haben wir die paarweise Unabhängigkeit der X_i benutzt, also

$$\text{Cov}(X_i, X_j) = \begin{cases} \sigma^2 & \text{für } i = j, \\ 0 & \text{für } i \neq j. \end{cases}$$

Als nächstes berechnen wir

$$\begin{aligned}
 E((X_k - \bar{X})^2) &= E(((X_k - \mu) - (\bar{X} - \mu))^2) \\
 &= E((X_k - \mu)^2 - 2(X_k - \mu)(\bar{X} - \mu) + (\bar{X} - \mu)^2) \\
 &= E((X_k - \mu)^2) - 2E((X_k - \mu)(\bar{X} - \mu)) + E((\bar{X} - \mu)^2) \\
 &= \text{Var}(X_k) - \frac{2}{n} \sum_{l=1}^n \underbrace{E((X_k - \mu)(X_l - \mu))}_{=\text{Cov}(X_k, X_l)} + \frac{1}{n} \sigma^2 \\
 &= \sigma^2 - \frac{2}{n} \sigma^2 + \frac{1}{n} \sigma^2 \\
 &= \frac{n-1}{n} \sigma^2.
 \end{aligned}$$

Nach diesen Vorbereitungen berechnen wir den Erwartungswert des Schätzers T_2 .

$$\begin{aligned}
 E(T_2) &= E\left(\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2\right) \\
 &= \frac{1}{n} \sum_{i=1}^n E((X_i - \bar{X})^2) \\
 &= \frac{1}{n} \cdot n \cdot \frac{n-1}{n} \sigma^2 \\
 &= \frac{n-1}{n} \sigma^2
 \end{aligned} \tag{10.6}$$

Der Schätzer T_2 ist also *nicht* erwartungstreu.

3. Aus (10.6) folgt sofort, dass

$$\begin{aligned}
 T_3 : \chi &\rightarrow \mathbb{R}, \\
 x = (x_1, \dots, x_n) &\mapsto \frac{1}{n-1} \sum_{i=1}^n (x_i - T_0(x))^2.
 \end{aligned}$$

(mit $n \geq 2$) ein erwartungstreuer Schätzer für die Varianz ist.

4. Im speziellen Falle des Münzwurfs (s. Beispiel 10.1.1) hängt das Verteilungsmaß der X_i und somit auch deren Varianz allein vom Parameter p ab. Man könnte auch, ausgehend von der Beziehung $\sigma^2 = p(1-p)$, einen Schätzer T_4 für die Varianz konstruieren:

$$\begin{aligned}
 T_4(x) &:= \bar{x}(1 - \bar{x}) \\
 &= \frac{1}{n} \sum_{i=1}^n x_i - \frac{1}{n^2} \sum_{k=1}^n \sum_{i=1}^n x_i x_k,
 \end{aligned}$$

wobei wir die Notation $\bar{x} = \frac{1}{n}(x_1 + \dots + x_n)$ verwendet haben. Aber auch dieser Schätzer ist nicht erwartungstreu:

$$\begin{aligned}
 E(T_4) &= E\left(\frac{1}{n} \sum_{k=1}^n X_k - \frac{1}{n^2} \sum_{k=1}^n \sum_{l=1}^n X_k \cdot X_l\right) \\
 &= \frac{1}{n} \cdot n \cdot p - \frac{1}{n^2} \sum_{k=1}^n E(X_k^2) - \frac{1}{n^2} \sum_{\substack{1 \leq k, l \leq n, \\ k \neq l}} E(X_k \cdot X_l) \\
 &= p - \frac{1}{n} p - \frac{n(n-1)}{n^2} p^2 \\
 &= \frac{n-1}{n} (p - p^2) \\
 &= \frac{n-1}{n} \sigma^2
 \end{aligned}$$

Allgemein können wir diese Beobachtung so formulieren: Ist $f : \mathbb{R} \rightarrow \mathbb{R}$ eine beliebige Funktion. Dann folgt aus der Erwartungstreue eines Schätzers T für einen reellen Parameter ϑ i.a. *nicht* die Erwartungstreue der Schätzers $f \circ T$ von $f(\vartheta)$.

Bemerkung 10.1.8. (Asymptotische Erwartungstreue und Konsistenz einer Folge von Schätzern)

1. In Beispiel 10.1.7 haben wir für jedes $n \geq 2$ die Schätzer $T_i^{(n)}$ mit $i \in \{1, 2, 3, 4\}$ definiert, von denen nur $T_1^{(n)}$ und $T_3^{(n)}$ erwartungstreu sind. Wir sehen aber auch, dass für „große“ n und für jede Realisierung $x = (x_1, \dots, x_n)$ die geschätzten Werte $T_i^{(n)}(x)$ nahe beieinander liegen, sich diese Schätzer bei praktischen Problemen mit „großem“ n nicht wesentlich voneinander unterscheiden. Diese Familien von Schätzern sind nämlich alle *asymptotisch erwartungstreu*.

Eine Familie $(T^{(n)})_{n \in \mathbb{N}}$ von Schätzern für einen Parameter $g(\vartheta)$ heißt **asymptotisch erwartungstreu**, wenn

$$\lim_{n \rightarrow \infty} E(T^{(n)}) = g(\vartheta).$$

2. Der geschätzte Wert $T_i^{(n)}(x^{(n)})$ hängt von speziellen Realisierungen $x^{(n)}$ ab. Man kann zeigen, dass die Familien $(T_i^{(n)})_{n \geq 2}$ **konsistent** sind, d.h. für jedes $\epsilon > 0$ gilt

$$\lim_{n \rightarrow \infty} P(\{x^{(n)} \in \chi_1^n : |T^{(n)}(x^{(n)}) - \sigma^2| > \epsilon\}) = 0.$$

Gleiches gilt für den Schätzer T_0 des Erwartungswertes (s. Beispiel 10.1.6). D.h. für festes $\epsilon > 0$ geht mit immer größer werdender Anzahl von Einzelexperimenten die Wahrscheinlichkeit dafür, dass der geschätzte Wert eines Parameters vom tatsächlichen Wert um mehr als ϵ abweicht, gegen Null. Man vergleiche dies mit dem schwachen Gesetz der großen Zahlen (Satz 9.1.55)

10.1.3 Konfidenzintervalle

Wir betrachten wieder ein n -fach wiederholtes Zufallsexperiment mit voneinander unabhängigen Einzelexperimenten. Diese seien durch voneinander unabhängige, identisch verteilte reelle Zufallsvariablen X_i mit Werten in χ_1 beschrieben. Ein Schätzer ordnet jeder Realisierung $(x_1, \dots, x_n) \in \chi_1^n \subset \mathbb{R}^n$ einen Schätzwert eines Parameters zu, dessen tatsächlicher Wert unbekannt ist. Für „große“ n liegt der Schätzwert mit großer Wahrscheinlichkeit nahe beim tatsächlichen Wert des Parameters, aber Abweichungen sind trotzdem möglich, wenn auch nur mit geringer Wahrscheinlichkeit. Z.B. kann bei 100-fachem Münzwurf mit einer fairen Münze 100-mal „Kopf“ geworfen werden, und in solchen seltenen Fällen wird der geschätzte Wert für den Erfolgsparameter der Münze vom tatsächlichen stark abweichen.

Wir möchten nun Aussagen über solche Abweichungen machen. Dazu geben wir zu jeder Realisierung $x = (x_1, \dots, x_n)$ nicht nur einen Schätzwert $\hat{\vartheta}$ an (den allgemeineren Fall, dass nicht ϑ , sondern $g(\vartheta)$ zu schätzen ist, beachten wir für den Augenblick nicht), sondern auch noch ein Intervall $[\hat{\vartheta}_1, \hat{\vartheta}_2] \ni \hat{\vartheta}$. Die Intervallgrenzen $\hat{\vartheta}_1$ und $\hat{\vartheta}_2$ sowie $\hat{\vartheta}$ können wir wieder als Zufallsvariablen betrachten, da sie Funktionen der zufälligen Werte $(x_1, \dots, x_n)$ sind, also $\hat{\vartheta}_1(x_1, \dots, x_n)$

etc. Das somit zufällige Intervall $[\hat{\vartheta}_1, \hat{\vartheta}_2]$ soll idealerweise mit großer Wahrscheinlichkeit den tatsächlichen Wert ϑ enthalten. Allerdings ist es auch wünschenswert, dass die Breite $|\hat{\vartheta}_2 - \hat{\vartheta}_1|$ möglichst klein ist. Diese Forderungen an das Zufallsintervall bestimmen z.B., wie groß n zu wählen ist, d.h. wie viele Einzelexperimente gemacht werden müssen.

Definition 10.1.9 (Konfidenzintervall).

Sei ein Schätzproblem (s. Definition 10.1.2) mit Stichprobenraum $\chi = \chi_1^n$ gegeben, und sei ϑ der zu schätzende Parameter. Seien

$$\hat{\vartheta}_i : \chi \rightarrow \mathbb{R}$$

(mit $i = 1, 2$) reelle Zufallsvariablen, also Funktionen, die jeder Realisierung $x = (x_1, \dots, x_n) \in \chi$ jeweils eine Zahl $\hat{\vartheta}_1(x)$, bzw. $\hat{\vartheta}_2(x)$ zuordnen.. Dann heißt das Zufallsintervall $[\hat{\vartheta}_1, \hat{\vartheta}_2]$ ein **Konfidenzintervall** oder auch **Vertrauensintervall** für den Parameter ϑ mit **Konfidenzniveau** $\gamma \in [0, 1]$, wenn

$$\forall \vartheta \in \Theta \quad P_\vartheta(\hat{\vartheta}_1 \leq \vartheta \leq \hat{\vartheta}_2) \geq \gamma$$

gilt, d.h.

$$\forall \vartheta \in \Theta \quad P_\vartheta(\{x \in \chi\} \mid \hat{\vartheta}_1(x) \leq \vartheta \leq \hat{\vartheta}_2(x)) \geq \gamma.$$

Bemerkung 10.1.10 (zum Konfidenzniveau).

In der Situation von Definition 10.1.9 wird jeder Realisierung $x = (x_1, \dots, x_n)$ ein von x abhängiges Intervall zugeordnet. Die Wahrscheinlichkeit (bzgl. des Maßes P_ϑ) der Menge derjenigen Realisierungen, die zu einem Intervall führen, das den tatsächlichen Wert ϑ enthält, soll mindestens γ betragen. Und dies muß für alle Maße P_ϑ gelten, die bei dem Schätzproblem betrachtet werden.

Die Angabe von Konfidenzintervallen ist im allgemeinen nicht einfach und hängt natürlich auch von der betrachteten Familie $(P_\vartheta)_{\vartheta \in \Theta}$ von Wahrscheinlichkeitsmaßen ab. Wir beschränken uns hier auf die Diskussion des einfachsten Falles, dem der Normalverteilung, der jedoch gemäß dem zentralen Grenzwertsatz zu vielen anderen Fällen eine brauchbare Approximation liefert.

Beispiel 10.1.11. (Konfidenzintervall für unabhängige $\mathcal{N}(\mu, \sigma^2)$ -verteilte Zufallsvariablen mit bekanntem σ^2 und zu schätzendem μ)

Seien $X_1, \dots, X_n$ voneinander unabhängige und $\mathcal{N}(\mu, \sigma^2)$ -verteilte Zufallsvariablen, die wir als zufällig gestreute Meßergebnisse interpretieren können. Sei σ^2 bekannt und sei μ z.B. mit einem Konfidenzniveau $\gamma = 0.95$ zu schätzen.

Mann kann zeigen, dass auch die Zufallsvariable $X = X_1 + \dots + X_n$ normalverteilt ist, und zwar mit Erwartungswert $n\mu$ und Varianz $n\sigma^2$. Somit ist auch die Zufallsvariable $\bar{X} = \frac{1}{n}(X_1 + \dots + X_n)$ normalverteilt, mit Erwartungswert μ und Varianz $\frac{1}{n}\sigma^2$. Des Weiteren ist

$$Z = \frac{\sqrt{n}(\bar{X} - \mu)}{\sigma}$$

$\mathcal{N}(0, 1)$ -verteilt. Wegen $\Phi(1.96) \approx 0.975$ (s. (9.49)) ist

$$\int_{-1.96}^{1.96} f_{0,1}(y) dy \approx 0.95,$$

also

$$P_{\mu, \sigma^2}(|Z| \leq 1.96) \approx 0.95.$$

Die Bedingung $|Z| \leq 1.96$ können wir umschreiben als

$$\begin{aligned} \frac{\sqrt{n}(\bar{X} - \mu)}{\sigma} &\leq 1.96 \\ \Leftrightarrow |\bar{X} - \mu| &\leq \frac{1.96 \cdot \sigma}{\sqrt{n}} \\ \Leftrightarrow |\bar{X} - \mu| &\leq \frac{2\sigma}{\sqrt{n}}. \end{aligned}$$

Wir verwenden nun $\bar{X}$ als (erwartungstreuen) Schätzer für μ . Mit einer Wahrscheinlichkeit von etwa 0.95 weicht dann der zufällige Schätzwert vom tatsächlichen Wert μ um höchstens $\frac{2\sigma}{\sqrt{n}}$ ab. Also ist das (zufällige) Intervall $[\bar{X} - \frac{2\sigma}{\sqrt{n}}, \bar{X} + \frac{2\sigma}{\sqrt{n}}]$ ein Konfidenzintervall zum Konfidenzniveau 0,95. D.h. die Wahrscheinlichkeit für eine Realisierung $x = (x_1, \dots, x_n)$, die zu einem Schätzwert $\bar{x} = \frac{1}{n}(x_1 + \dots + x_n)$ und einem Intervall

$$[\hat{\vartheta}_1(x), \hat{\vartheta}_2(x)] = [\bar{x} - \frac{2\sigma}{\sqrt{n}}, \bar{x} + \frac{2\sigma}{\sqrt{n}}] \quad (10.7)$$

führt, das den tatsächlichen Erwartungswert μ *nicht* enthält, ist nicht größer als 0.05. Die Breite des Intervalls ist proportional zu $\frac{1}{\sqrt{n}}$, wird also mit wachsendem n immer kleiner.

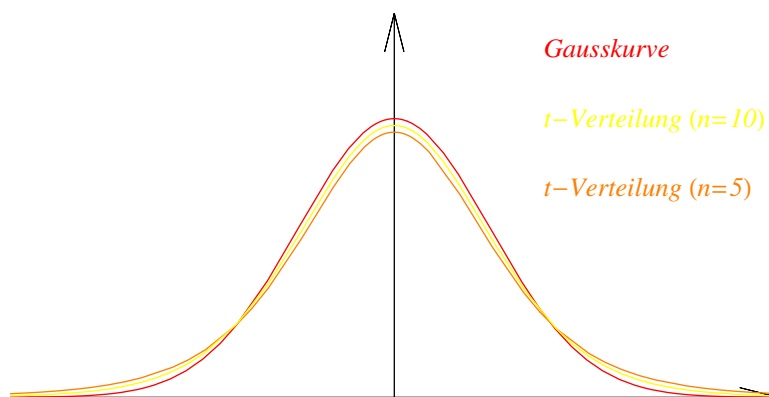


Abbildung 10.1: Die t - oder Student-Verteilung

Bemerkung 10.1.12. (Schätzung des Erwartungswertes bei unbekannter Varianz, t -Verteilung)

Wenn nun die Varianz auch unbekannt ist, muß auch sie geschätzt werden. Für große n kann man in guter Näherung in (10.7) die Streuung σ durch einen Schätzwert $\hat{\sigma}$ ersetzen.

Für kleine n benutzt man zur Konstruktion von Konfidenzintervallen die Quantile der so genannten t -Verteilung (oder auch Student-Verteilung), die wir in Abbildung 10.1 für $n = 5$ und $n = 10$

skizziert haben. Anschaulich gesprochen ist diese Verteilung für jedes $n < \infty$ etwas „breiter“ als die Normalverteilung, aber für große n geht sie rasch in die Normalverteilung über.

Für die Schätzung eines konkreten Erwartungswertes bedeutet dies, dass bei nur geschätzter Streuung zu gegebenem Konfidenzniveau die Konfidenzintervalle etwas größer ausfallen als in (10.7). Man muss sozusagen noch etwas mehr Sicherheit einplanen, da die Streuung nicht exakt bekannt ist.

Die t -Verteilung spielt auch eine wichtige Rolle beim sogenannten t -Test, den wir sehr kurz in Abschnitt 10.2.3 behandeln. Für weitere Informationen verweisen wir auf die Fachliteratur, z.B. auf [Bos00, Sac02, Sta02].

10.1.4 Empirischer Median einer Stichprobe

Wir wollen hier noch kurz eine wichtige Kenngröße zur Beschreibung einer Stichprobe erwähnen, den **empirischen Median**. Ähnlich wie der Durchschnitt einer Stichprobe ein Schätzer für den Erwartungswert einer Verteilung darstellt, stellt der empirische Median einen Schätzer für den Median (also das $\frac{1}{2}$ -Quantil) einer Verteilung dar.

Definition 10.1.13 (empirischer Median). Sei $x \in \mathbb{R}^n$ eine geordnete Stichprobe, $x_1 \leq x_2 \leq \dots \leq x_n$. Der **empirische Median** $\tilde{x}$ dieser Stichprobe ist definiert als

$$\tilde{x} := \begin{cases} x_{\frac{n+1}{2}} & \text{falls } n \text{ ungerade,} \\ \frac{1}{2} \left(x_{\frac{n}{2}} + x_{\frac{n+2}{2}} \right) & \text{falls } n \text{ gerade.} \end{cases}$$

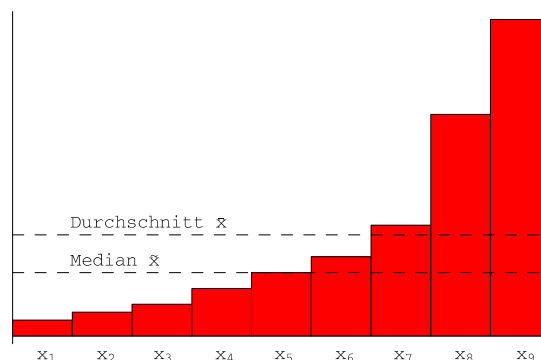


Abbildung 10.2: Durchschnitt und Median einer Stichprobe.

Wir diskutieren den Unterschied zwischen Durchschnitt und Median, der in Abbildung 10.2 illustriert ist, am leichtesten anhand eines Beispiels.

Beispiel 10.1.14 (Durchschnitt und Median).

Jahresgehälter von 5 zufällig ausgewählten Angestellten (in 1000 Euro)

$$\begin{aligned}x &= (22, 28, 40, 60, 850), & \left(\bar{x} = \frac{1000}{5} = 200\right) \\ \tilde{x} &= x_3 = 40,\end{aligned}$$

bzw. von 6 zufällig ausgewählten Angestellten:

$$\begin{aligned}x &= (22, 28, 40, 60, 60, 850), & \left(\bar{x} = \frac{1060}{6} = 210\right) \\ \tilde{x} &= \frac{x_3 + x_4}{2} = \frac{40 + 60}{2} = 50.\end{aligned}$$

Man sieht, dass der letzte Wert von 850, der wesentlich höher ist als die anderen Werte der Stichprobe (vielleicht das Einkommen eines CEO), den Median überhaupt nicht beeinflusst, den Durchschnittswert hingegen stark. Häufig wird der Median benutzt, wenn man einen Schätzer konstruieren will, der unempfindlich gegen Ausreißer ist.

10.2 Hypothesentest

Wir wollen in diesem Abschnitt bereits einmal kurz auf eine sehr wichtige Anwendung der Statistik eingehen, den Test von Hypothesen. Bevor wir die Problematik an einem einfachen Beispiel illustrieren und danach auf immer komplexere Fälle eingehen, möchten wir ganz kurz einige Begriffe wiederholen, die wir in diesem Abschnitt benötigen werden.

10.2.1 Hilfsmittel

Wir erinnern an die wichtigsten beiden Schätzer, die wir auch für den Hypothesentest benötigen. Sei $x = (x_1, \dots, x_n) \in \mathbb{R}^n$ eine Stichprobe vom Umfang n , mit Merkmalswerten $x_i \in \mathbb{R}$. Dann heisst

1. $\bar{x} = \frac{\sum_{i=1}^n x_i}{n}$ der (*empirische*) *Mittelwert*, und
2. $s^2 = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n-1}$ die (*empirische*) *Varianz* oder *mittlere quadratische Abweichung* der Stichprobe.

Hat jede der Zufallsvariablen X_i den Erwartungswert μ und die Varianz σ^2 , dann hat $\bar{X}$ den Mittelwert μ und die Varianz $\frac{\sigma^2}{n}$. Nach dem zentralen Grenzwertsatz ist $\bar{X}$ für große n sogar näherungsweise $\mathcal{N}(\mu, \frac{\sigma^2}{n})$ verteilt.

Die Verteilungsfunktion $\Phi : \mathbb{R} \rightarrow \mathbb{R}$ der $\mathcal{N}(0, 1)$ Verteilung ist definiert durch

$$\Phi(z) := \int_{-\infty}^z \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}} dx,$$

siehe auch Abbildung 9.12.

10.2.2 Ablehnungs- und Verträglichkeitsbereich

Die Frage, die man sich beim Hypothesentest stellt, ist immer: ist eine bestimmte Hypothese H_0 verträglich mit den experimentellen Tatsachen? In der Statistik wird es dafür immer nur Wahrscheinlichkeitsaussagen geben, d.h. man fragt sich: mit *welchem Konfidenzniveau* γ ist die Hypothese H_0 verträglich mit den experimentellen Tatsachen?

Meist führt man für den Hypothesentest eine sogenannte *Prüfgröße* ein, die durch das Experiment bestimmt wird, und die uns helfen soll, zu entscheiden, ob wir die Hypothese ablehnen müssen, oder ob wir sie akzeptieren können. Nehmen wir nun also an, jemand stellt die Hypothese H_0 auf:

$$H_0 : \quad \text{Die Zufallsvariable } X \text{ ist } \mathcal{N}(\mu, \sigma^2) \text{ verteilt.}$$

Diese Hypothese möchten wir testen. Um sie zu testen, nehmen wir eine Stichprobe $(x_1, \dots, x_n)$ vom Umfang n , und bilden den Mittelwert

$$\bar{x} := \frac{1}{n} \sum_{i=1}^n x_i,$$

der uns als Prüfgröße dienen soll. Wenn H_0 wahr ist, also X normalverteilt ist, ist auch die Zufallsvariable $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$ normalverteilt, mit Erwartungswert μ und Varianz $\frac{\sigma^2}{n}$. Ganz ähnlich, wie wir zuvor Konfidenzintervalle bestimmt haben, gehen wir jetzt vor. Wir bestimmen zunächst zu gegebenem Konfidenzniveau γ bzw. zu gegebener Fehlerwahrscheinlichkeit $\alpha = 1 - \gamma$ ein Quantil z_0 der Standard-Normalverteilung so dass

$$\Phi(z_0) = 1 - \frac{\alpha}{2} \tag{10.8}$$

(z.B. $z_0 = 1.96$ für $\alpha = 5\%$). Sodann unterscheiden wir:

- a) *Ablehnungsbereich*: Falls $|\bar{x} - \mu| \geq z_0 \frac{\sigma}{\sqrt{n}}$ lehnen wir die Hypothese H_0 ab. Die Irrtumswahrscheinlichkeit ist α : Wenn die Hypothese H_0 in Wirklichkeit wahr ist und wenn wir das gesamte Experiment viele Male wiederholen würden, dann träte das Ablehnungs-Ereignis $|\bar{x} - \mu| \geq z_0 \frac{\sigma}{\sqrt{n}}$ nur bei einem Anteil α der Experimente auf. Um dies zu sehen, betrachtet man

$$\begin{aligned} P\left(|\bar{X} - \mu| \geq z_0 \frac{\sigma}{\sqrt{n}}\right) &= P\left(\left|\underbrace{\frac{\sqrt{n}\bar{X} - \mu}{\sigma}}_{=:Z}\right| \geq z_0\right) \\ &= P(Z \leq -z_0 \text{ oder } z_0 \leq Z) \\ &= \frac{\alpha}{2} + \frac{\alpha}{2}, \end{aligned}$$

denn die Zufallsvariable Z ist $\mathcal{N}(0, 1)$ verteilt.

- b) *Nicht-Ablehnungsbereich*: Falls nun aber $|\bar{x} - \mu| < z_0 \frac{\sigma}{\sqrt{n}}$, dann können wir die Hypothese H_0 nicht mit Sicherheit γ ablehnen. Man sagt dann, die Hypothese sei mit dem Experiment *verträglich*. Allerdings können wir auch nicht behaupten, die Hypothese H_0 sei bewiesen, denn unsere Beobachtung wäre ebensogut (oder sogar besser) verträglich mit anderen Hypothesen (dass X beispielsweise $\mathcal{N}(\bar{x}, \sigma^2)$ verteilt sei). Frei nach Wittgensteins „Wovon man nicht sprechen kann, darüber muß man schweigen“ haben sich die Statistiker deshalb entschieden, in diesem Fall einfach die Hypothese als nicht widerlegt zu betrachten und nichts weiter in die Ergebnisse der Stichprobe hineinzuinterpretieren.

Wir fassen also nochmal zusammen, dass man die Hypothese H_0 zwar bei entsprechenden experimentellen Ergebnissen mit einer gewissen Konfidenz ablehnen kann, dass man sie aber nicht mit Hilfe des Experiments beweisen kann. Es ist interessant, diese Asymmetrie mit der grundsätzlichen Problematik naturwissenschaftlicher Erkenntnis zu vergleichen, auf die Philosophen wie David Hume oder später Karl Popper hingewiesen haben, dass nämlich naturwissenschaftliche Hypothesen durch Experimente eindeutig widerlegt (falsifiziert), aber nicht wirklich bestätigt (verifiziert) werden können.

Beispiel 10.2.1 (Molekulargewichtsmessung).

In einem Fachartikel wird die atomare Struktur eines bislang unentschlüsselten Makromoleküls angegeben, die ein Molekülgewicht von genau $\mu = 1294$ u impliziert. Wir haben eine kleine Probe und wollen die Hypothese mit einem Massenspektrographen testen. Der Spektrograph ermittelt Massen sehr genau, aber mit einem normalverteiltem Fehler von $\pm 1\%$, und um die Masse genauer zu ermitteln, führen wir $n = 100$ Massenmessungen durch. (**Achtung:** dies geht nur, wenn wir annehmen können, dass die Messfehler wirklich unabhängig voneinander sind, und nicht z.B. durch einen Fehler in der Eichung verursacht sind.) Als Ergebnis erhalten wir den Mittelwert $\bar{x} = 1298$ u $\neq 1294$ u. Sind die Ergebnisse des Fachartikels durch diese Diskrepanz widerlegt?

Die zu testende Hypothese wäre

$$H_0 : \quad \begin{array}{l} \text{Die Massenmessungen haben den Erwartungswert } \mu = 1294 \text{ u} \\ \text{und eine Standardabweichung von } \sigma = 1\% \mu = 12.9 \text{ u.} \end{array}$$

Wir geben uns ein Konfidenzniveau von $\gamma = 95\%$ vor (also eine Irrtumswahrscheinlichkeit $\alpha = 5\%$), und bestimmen das Quantil z_0 zu 1.96, denn $\Phi(1.96) = 97.5\% = 1 - \frac{\alpha}{2}$. Sodann berechnen wir

$$z_0 \frac{\sigma}{\sqrt{n}} = 1.96 \cdot \frac{12.9 \text{ u}}{\sqrt{100}} = 1.96 \cdot 1.29 \text{ u} = 2.53 \text{ u}$$

sowie die Abweichung $|\bar{x} - \mu| = |1294 \text{ u} - 1298 \text{ u}| = 4 \text{ u}$. Aus der Tatsache, dass $4 \text{ u} > 2.53 \text{ u}$ schliessen wir, dass wir mit einer Sicherheit von 95% davon ausgehen können, dass wir die Ergebnisse des Fachartikels widerlegt haben.

Hinweis für Interessierte: Tatsächlich können wir uns allerdings noch sicherer sein. Wie sicher, berechnen wir wie folgt: wir bilden den Quotienten

$$z := \sqrt{n} \frac{\bar{x} - \mu}{\sigma} = \frac{4 \text{ u}}{1.29 \text{ u}} = 3.10$$

und verwenden die Umkehrung von (10.8):

$$\Phi(z) = 1 - \frac{\alpha}{2} \Leftrightarrow \alpha = 2(1 - \Phi(z)) \approx 2(1 - 0.9986) = 0.28\%,$$

d.h. wir dürfen uns zu $\gamma = 1 - \alpha = 99.72\%$ sicher sein, die Ergebnisse widerlegt zu haben.

10.2.3 Der t -Test

Schwieriger wird es bei Hypothesen der Form:

H_0 : Die Zufallsvariable X ist normalverteilt und hat den Erwartungswert μ .

Hier ist die Schwierigkeit, dass über die Varianz σ der Verteilung von X nichts gesagt wird. Wir ziehen eine Stichprobe $(x_1, \dots, x_n)$ von der wir wieder den Mittelwert $\bar{x}$ bilden. Glücklicherweise können wir die Varianz der Variable X durch

$$s^2 = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n - 1}$$

schätzen. Wir wählen jetzt als Prüfgröße die Variable

$$t = \sqrt{n} \frac{\bar{x} - \mu}{s} \quad \text{bzw.} \quad T = \sqrt{n} \frac{\bar{X} - \mu}{S}.$$

Das Subtrahieren des als bekannt angenommenen Erwartungswertes μ und Teilen durch die empirische Varianz S hat zur Folge, dass die so erhaltene Zufallsvariable T , ganz unabhängig von der wahren Varianz σ , einer wohldefinierten Verteilung folgt, die man die t - oder Student-Verteilung nennt. Diese Verteilung hängt allerdings von der Größe n der Stichprobe ab; sie geht für $n \rightarrow \infty$ in die $\mathcal{N}(0, 1)$ Verteilung über, siehe auch Abbildung 10.1. Bei festem n können wir zu gegebenem Konfidenzniveau $\gamma = 1 - \alpha$ ein Quantil t_0 der entsprechenden t -Verteilung bestimmen, analog zu (10.8), wo z_0 als Quantil der $\mathcal{N}(0, 1)$ Verteilung bestimmt wurde. Sodann unterscheiden wir:

- a) Falls $|t| \geq t_0$ lehnen wir die Hypothese H_0 ab (mit Irrtumswahrscheinlichkeit α).
- b) Falls $|t| < t_0$ sagen wir, die Hypothese H_0 sei verträglich mit dem Experiment.

t -Test ohne Normalverteilungsannahme

Noch schwieriger wird es, wenn man nur Hypothesen der Form

H_0 : Die Zufallsvariable X hat den Erwartungswert μ .

testen will, bei denen keine Annahme über die Art der Verteilung gemacht wird. Hier hilft der zentrale Grenzwertsatz, der besagt, dass der empirische Mittelwert, als Zufallsvariable $\bar{X}$ aufgefasst, für große n immer normalverteilt ist, ganz egal wie die ursprüngliche Verteilung war, mit gleichem Erwartungswert und einer um den Faktor n reduzierten Varianz. Wenn n groß ist, dürfen wir hier also auch den t -Test verwenden. Da n groß sein muss, machen wir aber auch keinen großen Fehler, wenn wir direkt mit Quantilen der Normalverteilung arbeiten.

10.2.4 Test auf Häufigkeiten

Als ein weiteres interessantes Problem des Hypothesentests betrachten wir nun Stichproben, in denen die Zufallsvariable X nur die Werte 0 oder 1 annehmen kann. Dies könnte z.B. die Antwort auf die Frage sein, ob eine zufällig ausgewählte Person Linkshänder ist, 1, oder nicht, 0. Eine entsprechende zu testende Hypothese wäre dann z.B., dass der Anteil von Linkshändern in der Bevölkerung gerade μ beträgt.

Die Hypothese, die wir testen möchten, ist jetzt wieder von der Form

$$H_0 : \quad \text{Die Zufallsvariable } X \text{ hat den Erwartungswert } \mu,$$

aber wir können nun ausnutzen, dass X nur Werte $\in \{0, 1\}$ annimmt. Wir nehmen wieder eine Stichprobe und nehmen als Prüfgröße den Mittelwert $\bar{x}$, der jetzt auch als Anteil der positiven Individuen in der Stichprobe interpretiert werden kann. Wenn die Hypothese H_0 richtig ist, dann ist die Summe $\sum_{i=1}^n X_i$ eine binomialverteilte Zufallsvariable mit Erwartungswert $n\mu$ und Varianz $n\mu(1 - \mu)$. Da wir damit auch die Verteilung der Prüfgröße $\bar{X}$ kennen, könnten wir im Prinzip jetzt schon entsprechende Quantile der Binomialverteilung finden, um unseren Ablehnungs- und Verträglichkeitsbereich zu definieren. Wenn wir aber keine Quantile der Binomialverteilung berechnen wollen, können wir die Tatsache ausnutzen, dass nach dem zentralen Grenzwertsatz für große n die Verteilung von $\bar{X}$ in eine entsprechende Normalverteilung übergeht, mit Erwartungswert μ und Varianz $\frac{\mu(1-\mu)}{n}$. Dies erlaubt uns ganz genau wie zuvor in (10.8), zu gegebener Irrtumswahrscheinlichkeit α ein Quantil z_0 der $\mathcal{N}(0, 1)$ Verteilung mit $\Phi(z_0) = 1 - \frac{\alpha}{2}$ zu bestimmen, und dann wieder zu unterscheiden:

- a) Falls $|\bar{x} - \mu| \geq z_0 \sqrt{\frac{\mu(1-\mu)}{n}}$ lehnen wir H_0 ab.
- b) Falls $|\bar{x} - \mu| < z_0 \sqrt{\frac{\mu(1-\mu)}{n}}$ lehnen wir H_0 nicht ab.

Beispiel 10.2.2. Ein Forschungsteam hat bei einem großangelegten Screening mit einer neuen Untersuchungsmethode 5000 Personen untersucht und unter anderem herausbekommen, dass ein bestimmtes Allel bei 1405 der untersuchten Personen vorkam. Aus hier nicht näher genannten Gründen bezweifeln wir die Validität der benutzten Untersuchungsmethode, und wollen das genannte Teilergebnis einem Test unterziehen. Wir wählen zufällig 100 aus den 5000 noch konservierten Blutproben, und untersuchen sie mit einem weltweit anerkannten, aber wesentlich aufwendigeren Verfahren für die Existenz des Allels, das bei 23 Proben ein positives Ergebnis erzielt. Wir vergleichen diesen Anteil von 23% mit dem erwarteten Anteil von $\mu = \frac{1405}{5000} = 28.1\%$. Ist dieser Unterschied ausreichend, um das Ergebnis der neuen Methode anzuzweifeln? Die zu testende Hypothese ist

$$H_0 : \quad \text{Der Anteil von Trägern des Allels ist } \mu = 28.1\%.$$

Wir geben uns zunächst $z_0 = 2$, also ein Konfidenzniveau von 95.44 % vor. Sodann berechnen wir

$$z_0 \sqrt{\frac{\mu(1-\mu)}{n}} = 2 \sqrt{\frac{0.281(1-0.281)}{100}} \approx 2 \sqrt{20.20 \cdot 10^{-4}} \approx 2 \cdot 4.49 \cdot 10^{-2} = 8.98\%,$$

und dann die Abweichung $|\bar{x} - \mu| = |23\% - 28.1\%| = 5.1\%$. Diese Abweichung ist nicht signifikant, und unser Experiment steht nicht im Widerspruch zur Hypothese – wir müssen also nicht an der Validität der neuen Untersuchungsmethode zweifeln. Wir haben jedoch auch keinen Beweis für ihre Validität erhalten, denn dafür war der Test nicht angelegt.

10.2.5 Test auf Einhaltung eines Grenzwerts

Eine weiterer interessanter Fall tritt auf, wenn wir eine Hypothese der Form

H_0 : Die Zufallsvariable X hat einen Erwartungswert **kleiner** bzw. **größer als** μ_0

testen wollen, wobei wir im konkreten Fall evtl. noch Annahmen über die Art der Verteilung machen dürfen. Das Problem bei dem Test auf Einhaltung eines Grenzwertes ist, dass die Hypothese nichts über den wirklichen Erwartungswert μ von X sagt, mit dessen Hilfe wir bisher immer eine wohldefinierte Verteilung für unsere Prüfgröße definieren konnten, die uns erlaubte, die Irrtumswahrscheinlichkeit anzugeben. Jetzt wird nur eine Grenze postuliert, $\mu \leq \mu_0$ bzw. $\mu \geq \mu_0$.

Wir werden im folgenden ein Beispiel zur Motivation betrachten, das wie im vorherigen Unterabschnitt von Häufigkeiten handelt.

Beispiel 10.2.3 (Der Corn-Tester).

Ein Getreidehändler will eine Schiffsladung US-Futtermais kaufen. Es gab gerade ein Problem mit einer neuen Sorte genetisch manipulierten Maises, die schwach giftig für Schweine ist und deshalb nicht mehr verkauft werden darf. Getreideverkäufer versuchen häufig durch „*blending*“, also durch Mischen verschiedener Sorten Mais, das verbotene Material noch loszuwerden. Davor will sich der Getreidehändler schützen: Vor dem Kauf beauftragt er deshalb einen „*Corn-Tester*“, die Schiffsladung zu untersuchen. Wenn weniger als 1% der Körner von der problematischen Sorte sind, darf der Mais weiter verfüttert werden. Der Corn-Tester hat ein Testverfahren, mit dem er eindeutig feststellen kann, ob ein Mais Korn „schlecht“ ist, oder nicht, und er beherrscht die Regeln der Statistik. Er hat als Ziel, mit Irrtumswahrscheinlichkeit $\alpha = 5\%$ zu garantieren, dass sich in der Schiffsladung weniger als $\mu_0 = 1\%$ schlechte Körner befinden.

Er zieht eine Stichprobe von $n = 1000$ zufällig gewählten Körnern und untersucht sie, und erhält darin einen Anteil von $\bar{x}$ schlechten Körnern. Wenn er mehr als $\mu_0 n = 10$ schlechte Körner darunter findet, wird er sicher vom Kauf abraten. Wenn er weniger schlechte Körner findet, kann dies entweder daran liegen, dass die Schiffsladung in Ordnung ist, oder an einer zufälligen Schwankung der Zufallsvariablen $\bar{X}$. Was tun, um zu $\gamma = 1 - \alpha = 95\%$ garantieren zu können, dass die Schiffsladung weniger als $\mu_0 = 1\%$ schlechte Körner enthält? Wo sollte er die Grenze ziehen? Wir haben in den vorherigen Hypothesentests schon gesehen, dass man Sicherheit nur beim Ablehnen einer Hypothese haben kann. Der Corn-Tester hat deshalb als Ziel, die Hypothese

H_0 : Der Anteil schlechter Körner in der Ladung, μ , ist größer als μ_0 .

mit gegebener Irrtumswahrscheinlichkeit zu widerlegen.

Test der Hypothese $\mu > \mu_0$ beim Test auf Häufigkeiten

Im folgenden untersuchen wir genau diesen Fall. X sei eine Zufallsvariable, die die Werte 0 und 1 annimmt, und wir untersuchen die Hypothese

$$H_0 : \quad \text{Die Zufallsvariable } X \text{ hat einen Erwartungswert } \mu \geq \mu_0.$$

Als Prüfgröße wählen wir, wie in Abschnitt 10.2.4, den Mittelwert $\bar{x}$ der Stichprobe, der als Anteil der positiven Testergebnisse aufgefasst werden kann.

Wenn wir μ kennen würden, dann wüssten wir, wie zuvor, dass der Mittelwert $\bar{X}$ den Erwartungswert μ und die Varianz $\sqrt{\frac{\mu(1-\mu)}{n}}$ hätte, und dass er bei großem n normalverteilt ist. Wir geben uns wieder ein Konfidenzniveau $\gamma = 1 - \alpha$ vor, und suchen nun eine Zahl $c > 0$, so dass wir die Hypothese H_0 , also $\mu \geq \mu_0$, für $\bar{x} \leq \mu_0 - c$ sicher ablehnen können. Wir definieren uns also:

a) *Ablehnungsbereich:* $\bar{x} \leq \mu_0 - c$

b) *Nicht-Ablehnungsbereich:* $\bar{x} > \mu_0 - c$.

Die Schwierigkeit besteht darin, das richtige c zu einer gegebenen Irrtumswahrscheinlichkeit α zu finden. In Wahrscheinlichkeiten ausgedrückt, wollen wir sicher sein, so dass für **jeden** wahren Erwartungswert μ , der mit der Hypothese H_0 verträglich ist, also für jedes $\mu \geq \mu_0$, das Ablehnungsereignis $\bar{x} \leq \mu_0 - c$ nur mit Wahrscheinlichkeit $P_\mu(\bar{x} \leq \mu_0 - c)$ kleiner als α auftritt:

$$\forall \mu \geq \mu_0 : \quad P_\mu(\bar{x} \leq \mu_0 - c) \leq \alpha \quad \Leftrightarrow \quad \max_{\mu \geq \mu_0} P_\mu(\bar{x} \leq \mu_0 - c) \leq \alpha.$$

Wir können jetzt ausnutzen, dass die Variable $\bar{x}$ bei gegebenem μ (und großem n) eine $\mathcal{N}(\mu, \frac{\mu(1-\mu)}{n})$ verteilte Variable ist. Man kann nämlich zeigen (vgl. Abbildung 10.3), dass

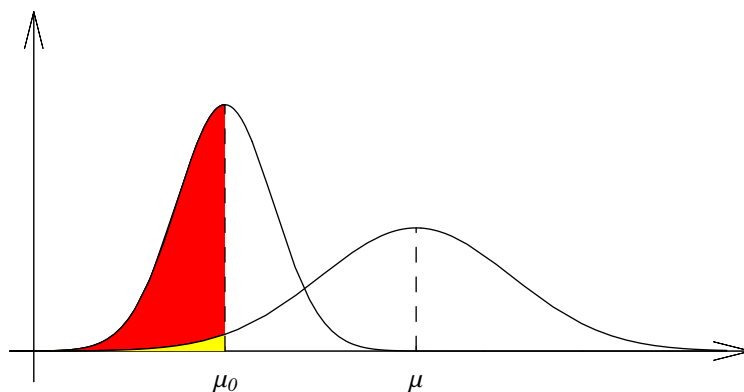


Abbildung 10.3: Die Irrtumswahrscheinlichkeit für $\mu = \mu_0$ ist größer als für $\mu > \mu_0$.

$$\max_{\mu \geq \mu_0} P_\mu(\bar{x} \leq \mu_0 - c) = P_{\mu_0}(\bar{x} \leq \mu_0 - c).$$

Daraus folgt, dass wir c aus der Gleichung

$$P_{\mu_0}(\bar{x} \leq \mu_0 - c) = \alpha$$

berechnen können, bzw. umgeformt

$$P_{\mu_0} \left(\underbrace{\sqrt{n} \frac{\bar{x} - \mu_0}{\mu_0(1 - \mu_0)}}_{\mathcal{N}(0,1) \text{ verteilt}} \leq \underbrace{\sqrt{n} \frac{-c}{\mu_0(1 - \mu_0)}}_{=:-z_0} \right) = \alpha \quad \Leftrightarrow \quad \Phi(-z_0) = \alpha.$$

So können wir wieder z_0 als ein Quantil der Normalverteilung zu gegebener Irrtumswahrscheinlichkeit α bestimmen. Man kann z_0 wegen der Symmetrie der Normalverteilung allerdings ebenso durch

$$\Phi(z_0) = \gamma = 1 - \alpha$$

ermitteln. Dann können wir unterscheiden:

- a) Falls $\bar{x} \leq \mu_0 - z_0 \sqrt{\frac{\mu_0(1-\mu_0)}{n}}$ lehnen wir H_0 ab (mit Irrtumswahrscheinlichkeit α)
- b) Falls $\bar{x} > \mu_0 - z_0 \sqrt{\frac{\mu_0(1-\mu_0)}{n}}$ lehnen wir H_0 nicht ab.

Beispiel 10.2.4 (Anwendung auf das Corn-Tester Beispiel).

Beim Corn-Tester Beispiel war $\mu_0 = 0.01$ und $n = 1000$, also $\sqrt{\frac{\mu_0(1-\mu_0)}{n}} \approx \sqrt{\frac{0.01}{1000}} \approx 0.003$. Wegen $\Phi(-1.64) = 5\%$ bzw. $\Phi(1.64) = 95\%$ setzen wir $z_0 = 1.64$ und berechnen die Grenze zu

$$\mu_0 - z_0 \sqrt{\frac{\mu_0(1-\mu_0)}{n}} = 0.01 - 1.64 \cdot 0.003 \approx 0.01 - 0.0049 \approx 0.5\% = \frac{5}{1000}.$$

Ist also $\bar{x} \leq 0.5\%$, dann ist der Corn-Tester zu $\gamma = 95\%$ sicher, dass es nicht mehr als $\mu_0 = 1\%$ schlechte Körner in der Gesamtladung gibt. Unter den 1000 getesteten Maiskörnern dürfen sich also maximal 5 Körner der problematischen Art befinden.

Test der Hypothese $\mu \geq \mu_0$ bei bekannter Varianz

Der Test auf Einhaltung eines Grenzwertes wird wesentlich einfacher, wenn wir es mit einer kontinuierlichen Zufallsvariablen mit als bekannt angenommener Varianz σ zu tun haben. Wir untersuchen die Hypothese

$$H_0 : \quad \text{Die Zufallsvariable } X \text{ ist } \mathcal{N}(\mu, \sigma^2) \text{ verteilt, mit } \mu \geq \mu_0.$$

Als Prüfgröße wählen wir wieder den Mittelwert $\bar{x}$, und zu gegebenem Konfidenzniveau γ bestimmen wir z_0 so dass $\Phi(z_0) = \gamma$. Dann können wir unterscheiden:

- a) Falls $\bar{x} \leq \mu_0 - z_0 \frac{\sigma}{\sqrt{n}}$ lehnen wir H_0 ab (wir sind also mit Irrtumswahrscheinlichkeit $\alpha = 1 - \gamma$ sicher, dass $\mu < \mu_0$).
- b) Falls $\bar{x} > \mu_0 - z_0 \frac{\sigma}{\sqrt{n}}$ lehnen wir H_0 nicht ab.

Test der Hypothese $\mu \leq \mu_0$ bei bekannter Varianz

Umgekehrt gilt natürlich für den Test der Hypothese

H_0 : Die Zufallsvariable X ist $\mathcal{N}(\mu, \sigma^2)$ verteilt, mit $\mu \leq \mu_0$:

- a) Falls $\bar{x} \geq \mu_0 + z_0 \frac{\sigma}{\sqrt{n}}$ lehnen wir H_0 ab (wir sind also mit Irrtumswahrscheinlichkeit α sicher, dass $\mu > \mu_0$).
- b) Falls $\bar{x} < \mu_0 + z_0 \frac{\sigma}{\sqrt{n}}$ lehnen wir H_0 nicht ab.

Beispiel 10.2.5 (Ozon-Messung).

Ein Analysegerät misst die Ozonkonzentration μ in Mikrogramm je Kubikmeter Luft mit einer Standardabweichung von $\pm 10 \frac{\mu\text{g}}{\text{m}^3}$. Wir mitteln über $n = 4$ Messungen, und sollen Smogalarm geben, wenn der Grenzwert von $\mu_0 = 120$ Mikrogramm je Kubikmeter Luft mit Konfidenzniveau $\gamma = 95\%$ überschritten wird. Wie hoch muss der Mittelwert $\bar{x}$ sein, damit wir Alarm schlagen? Wir berechnen $z_0 = 1.64$ und $\frac{\sigma}{\sqrt{n}} = 5 \frac{\mu\text{g}}{\text{m}^3}$, und erhalten als Anweisung:

$$\text{Falls } \bar{x} \geq \mu_0 + z_0 \frac{\sigma}{\sqrt{n}} = (120 + 1.64 \cdot 5) \frac{\mu\text{g}}{\text{m}^3} = 128.2 \frac{\mu\text{g}}{\text{m}^3}$$

schlagen wir Alarm. Die Wahrscheinlichkeit eines Fehlalarms ist dabei 5 %.

Literaturverzeichnis

- [AE99] H. Amann and J. Escher. *Analysis I*. Birkhäuser, 1999.
- [Ama83] Herbert Amann. *Gewöhnliche Differentialgleichungen*. de Gruyter, Berlin; New York, 1983.
- [Bat80] Eduard Batschelet. *Einführung in die Mathematik für Biologen*. Springer, 1980.
- [BBDH01] Hans-Peter Beck-Bornholdt, Hans-Herrmann Dubben, and Imke Hoffmann. *Der Hund, der Eier legt. Erkennen von Fehlinformation durch Querdenken*. Rowohlt Taschenbuch, 2 edition, 2001.
- [BF] Martin Barner and Friedrich Flohr. *Analysis I*. de Gruyter.
- [Bos99] Karl Bosch. *Elementare Einführung in die Wahrscheinlichkeitsrechnung*. Vieweg, 7 edition, 1999.
- [Bos00] Karl Bosch. *Elementare Einführung in die angewandte Statistik*. Vieweg, 7 edition, 2000.
- [BSMM00] Ilja N. Bronstein, Konstantin A. Semendjajew, Gerhard Musiol, and Heiner Mühlig. *Taschenbuch der Mathematik*. Harri Deutsch Verlag, 2000.
- [Cre79] Hubert Cremer. *Carmina Mathematica und andere poetische Jugendsünden*. Verlag J.A. Mayer, Aachen, 6 edition, 1979.
- [Fis00] Gerd Fischer. *Lineare Algebra*. Vieweg Studium, 12 edition, 2000.
- [FK90] H. Fischer and H. Kaul. *Mathematik für Physiker. Band 1: Grundkurs*. Teubner, 2 edition, 1990.
- [FLS63] Richard P. Feynman, Robert B. Leighton, and Matthew Sands. *The Feynman Lectures on Physics, vol I*. Addison-Wesley Pub Co, 1963.
- [Fora] Forster. *Analysis I*. Vieweg.
- [Forb] Forster. *Analysis II*. Vieweg.

- [Haa10] A. Haar. Zur Theorie der orthogonalen Funktionen-Systeme. *Math. Ann.*, 69:331–371, 1910.
- [Jäh98] Klaus Jähnich. *Lineare Algebra*. Springer-Verlag, 4 edition, 1998.
- [Krä00] Walter Krämer. *So lügt man mit Statistik*. Piper, 2000.
- [Kre02] Ulrich Krengel. *Einführung in die Wahrscheinlichkeitstheorie und Statistik*. Vieweg, 6 edition, 2002.
- [Lip99] Seymour Lipschutz. *Lineare Algebra. Schaum's Überblicke und Aufgaben*. McGraw-Hill Germany/Hanser Fachbuchverlag, 2 edition, 1999.
- [Mur02] J.D. Murray. *Mathematical Biology*. Springer, 3 edition, 2002. ISBN: 0387952233.
- [Pap] Lothar Papula. *Mathematik für Ingenieure und Naturwissenschaftler*. Vieweg.
- [Sac02] Lothar Sachs. *Angewandte Statistik*. Springer, 10 edition, 2002.
- [Sch] Harald Scheid. *Folgen und Funktionen: Einführung in die Analysis*. Mathematische Texte. Spektrum.
- [Seg84] Lee A. Segel. *Modeling Dynamic Phenomena in Molecular and Cellular Biology*. Cambridge University Press, 1984. ISBN: 052127477X.
- [SG] H. Stoppel and B. Griesse. *Übungsbuch zur Linearen Algebra*. Vieweg.
- [SH] S. L. Salas and Einar Hille. *Calculus*. Spektrum.
- [SR94] Robert R. Sokal and F. James Rohlf. *Biometry*. W H Freeman & Co, 3 edition, 1994.
- [Sta02] Werner A. Stahel. *Statistische Datenanalyse*. Vieweg, 4 edition, 2002.
- [Str00] S. H. Strogatz. *Nonlinear Dynamics and Chaos*. Perseus Publishing, 2000. ISBN: 07328204536.
- [Vog94] Herbert Vogt. *Grundkurs Mathematik für Biologen*. Teubner, 1994.
- [Wal93] Wolfgang Walter. *Gewöhnliche Differentialgleichungen*. Springer, 1993. ISBN: 038756294X.
- [YHYS96] Edward K. Yeagers, James V. Herod, R. Yeagers, and R. Shonkwiler. *An Introduction to the Mathematics of Biology, With Computer Algebra Models*. Springer, 1996. ISBN: 0817638091.

Index

- Ähnlichkeit
 - von Matrizen, 152
- Ähnlichkeit von Matrizen, 152
- Abbildung, **35**
 - Injektiv, 36
 - surjektiv, 36
 - bijektiv, 36
 - Verknüpfung, 37
- Ableitung, 98, 271
 - Approximationseigenschaft der, **278**, 282
 - einer Kurve, 271
 - höhere, 99
 - partielle, 274, **277**
 - partielle höherer Ordnung, 281
 - totale, **278**, 282
- absolute Konvergenz, 84
- Achilles und die Schildkröte, 86
- affiner Unterraum, 66
- Alternierende Reihen, 83
- Ameisenhaufen, 288
- Analysis
 - im R^n , 269
- Anfangsbedingung
 - für die Wärmeleitungsgleichung, 200
- Anfangswert, 304
- Anfangswertproblem, 304
- Anfangswertproblem
 - partielle Differentialgleichung, 200
- Anfangszustand, 304
- Approximation
 - durch Ableitung, 278
- Aussage, **19**
 - äquivalent, 19
 - Disjunktion, 20
 - Konjunktion, 20
 - Negation, 20
 - Verneinung, 20
- Axiom, **25**
- Baumdiagramm, 212
- Bayes
 - Formel von, 210
- bedingte Konvergenz, 84
- bedingte Wahrscheinlichkeit, 207
- bedingte Wahrscheinlichkeit
 - Definition, 209
- Bernoulli-Experiment , 217
- Bernoulli-Ungleichung, 28
- Bernoulli-Verteilung, 217
- Bienaymé
 - Formel von, 225
- bijektiv, 36
- Bild, 36, 53
- Bilinearform, 184
- Binomialverteilung, 217
- Binomialverteilung
 - Erwartungswert, 223
 - Stabdiagramm, 218
- Binomischer Lehrsatz, 29
- Bogenlänge, 271
- Bolzano-Weierstraß
 - Satz von, 78
- Cauchy-Folge, 75
- Cauchy-Schwarz-Ungleichung, 46, 180
- charakteristische Gleichung, 143
- charakteristische Polynom, 143
- Cauchy-Kriterium, 83
- Deduktion, 26
- Definition, 25

- Determinante, **132**
 - Berechnung, 135
 - Eigenschaften, 132
- diagnostischer Test, 212
- diagnostischer Test
 - Effizienz, 213
 - Sensitivität, 214
- Diagobalisierbarkeit
 - hermitescher Matrizen, 198
- Diagonalisierbarkeit, 153
 - symmetrischer Matrizen, 196
- Dichtefunktion, 237
- Differentialgleichung
 - gewöhnliche, 304
- Differentialgleichung
 - partielle, 200
- Differentialrechnung
 - Fundamentalsatz der Differential- und Integralrechnung, 165
 - Zusammenhang mit Integralrechnung, 163
- Differentiationsregeln
 - Produktregel, 102
 - Kettenregel, 102
- differenzierbar
 - stetig, 278
- Differenzierbarkeit
 - fett, 97
- Dimension, 50
- Distanz
 - von Vektoren, 45
- Divergenz
 - Definition, 301
 - einer Folge, 76
- Dreiecksungleichung, 185
- dynamische Systeme
 - lineare, 314
- Dynamisches System, 303
- dynamisches System, 304
 - autonomes, 305
- e, siehe* Eulersche Zahl
- Effizienz, 213
- Eigenraum, 142
- Eigenvektor, 142
- Eigenwert, 142
- Einheitsmatrix, 56
- elektrisches Feld, 297
- Elementarereignis, 203, 204
- Elementarmatrix, 63
- Endomorphismus, 127
- Energie, 299
- Ereignis, 204
- Ereignis
 - sicheres, 204
 - unmögliches, 204
- Ergebnismenge, 204
- Ergebnisraum, 204
- Erwartungstreue, 252
- Erwartungstreue
 - asymptotische, 256
- Erwartungswert, 222
- Erwartungswert
 - der Binomialverteilung, 223
 - Eigenschaften, 223
 - einer reellen kontinuierlichen Wahrscheinlichkeitsverteilung, 239
 - einer reellen kontinuierlichen Zufallsvariablen, 238
 - einer reellen Zufallsvariablen, 222
 - einer vektorwertigen Zufallsvariablen, 222
 - einer Verteilung, 222
 - eines Schätzers, 252
- euklidisch
 - Norm, 184
 - Norm
 - in $\mathbb{R}^3$, 177
 - in $\mathbb{R}^n$, 177
- euklidische Norm, 45
- Eulersche Zahl, 76
- Exponentialfunktion, 86
 - Eigenschaften, 87
- Exponentialverteilung, 241
- Extrema, 95, 114
 - Hinreichendes Kriterium, 114
- fair

- Würfel, 203
- Federpendel, 308, 309
- Feinheit
 - einer Zerlegung, 161
- Feld, *siehe* Vektorfeld
- Fixpunkt
 - eines dynamischen Systems, 322
 - Stabilität eines, 322
- Folge, **71**
 - Cauchy-Folge, 75
 - Divergenz, 76
 - Grenzwert, 73
 - Häufungspunkt, 77
 - Infimum, 80
 - Konvergenz, 73
 - Limes inferior, 80
 - Limes superior, 80
 - monoton, 75
 - Nullfolge, 72
 - Schranke, 80
 - Supremum, 80
- Fourier-Entwicklung, 185
- Fourier-Entwicklung
 - Koeffizienten, 188
- Fourier-Koeffizienten, 188
- Fourier-Reihe, 189
- Fourier-Reihe
 - Anwendung, 192
 - Beispiel, 189
 - Konvergenz, 189
- Fundamentalsatz
 - der Algebra, 126
- Fundamentalsatz
 - der Differential- und Integralrechnung, 165
- Funktion
 - Differenzierbarkeit, 98
 - inverse, 97
 - Konkavität, 116
 - Konvexität, 116
 - Maximum, 106
 - mehrerer Argumente, 274
 - Minimum, 106
- Gasgesetz
 - Isobare, 280
- Gasgesetz, ideales, 279
- Gauss-Glocke, 290
- gemeinsame Verteilung, 226
- Gesetz der großen Zahlen
 - schwaches, 232, 233
 - starkes, 234
- Gleichverteilung
 - auf einem beschränkten Intervall, 239
 - auf einem endlichem Wahrscheinlichkeitsraum, 205
- Goldener Schnitt, 140
- Gradient, 280
- Gradientenfeld, 298
- Graph, 211
 - einer Abbildung, 35
- Gravitationsfeld, 297
- Grenzwert
 - einer Folge, 73
 - einer Funktion, 91
- Gruppe, **41**
 - Gruppenaxiome, 41
 - inverses Element, 41
 - neutrales Element, 41
- Häufungspunkt, 77
- harmonischer Oszillator, 308
- Helix, 273
 - Länge einer, 273
- hermitesch, 198
- hermitesch
 - Operator, 202
- Hesse-Matrix, 295
- Hinreichende Bedingung, 22, 295
- Homogenität
 - Norm, 185
- Hypothese, 25
- Identität, 37
- Imaginärteil, 119
- Indirekter Beweis, 26
- Induktion, 26

- Injektiv, 36
- Insulin-Zucker-Modell, 303, 304, 304, 323
- Integral, 162
- Integral
 - uneigentliches, 171
- Integralrechnung
 - Fundamentalsatz der Differential- und Integralrechnung, 165
 - Zusammenhang mit Differentialrechnung, 163
- Integration
 - im $\mathbb{R}^n$, 285
 - auf gekrümmten Gebiet, 286
 - in Kugelkoordinaten, 292
 - in Polarkoordinaten, 288
 - sukzessive, 285
- Integrationsregeln, 166
- Integrationsregeln
 - für uneigentliche Integrale, 174
 - partielle Integration, 166
 - Substitutionsregel, 167
- Integrierbarkeit, 162
- Inverse Matrix, 58
 - Berechnung, 59
- Isobare, 280
- Isomorphismus, 59
- Jacobi-Matrix, 282
- Körper, **42**
 - Distributivgesetz, 42
- Körperaxiome, 42
- kartesisches Produkt, 34, 35
- Kern, 53
- Kettenregel
 - für Jacobi-Matrizen
 - fett, 284
- Koeffizientenmatrix, 67
- komplexe Konjugation, 121, 122
- komplexe Zahl, **119**
 - Rechenregeln, 121
- Konfidenzintervall, 257
- Konfidenzintervall
 - für Normalverteilungen, 257
- Konfidenzniveau, 257
- Konvergenz
 - bedingte, 84
 - einer Folge, 73
- Koordinaten
 - sphärische, 292
- Koordinatentransformation
 - für lineare Abbildungen, 151
 - lineare für Matrizen, 150
 - lineare für Vektoren, 145, 149
- Korollar, 25
- Korrelation
 - bei Merkmalsverteilung, 227
 - Fehlinterpretation, 229
 - Interpretation, 229
 - Rechenbeispiel, 227
 - und Kausalität, 229
- Korrelationskoeffizient
 - Definition, 224
- Kovarianz, 222
- Kovarianz
 - Definition, 224
 - Eigenschaften, 225
- Konvergenz
 - absolute, 84
- Kugelkoordinaten, 292, 293
- Kurve, 270
- Kurvenintegral, 298
 - vektorielles, *siehe* Kurvenintegral
- Kurvenlänge, 271
 - Definition, 272
- Laplace-Operator, 198
- Laplacescher Entwicklungssatz, 136
- Lemma, 25
- Limes, 277
- lineare Abbildung, **51**
- lineare Abhängigkeit, 49
- lineare dynamische Systeme, 314
- Lineare Unabhängigkeit, 49
- lineares Gleichungssystem, 60
 - homogenes, 60

- inhomogen, 65
- Koeffizientenmatrix, 61
- Lösungsverfahren, 61
- Zeilenumformung, 63
- Linearkombination, 48
- Logarithmus, 86
 - Eigenschaften, 89
- Lotto, 206
- Münzwurf
 - n -facher, 217
- Majorante, 73
- Matrix, **54**
 - Addition, 55
 - Inversion, 58
 - Matrizenmultiplikation, 55
 - orthogonale, 192
 - Rechenregeln, 56
 - Regularität, 58
 - Skalarmultiplikation, 55
 - symmetrische, 195
 - transponierte, 55
- Median, 248, 259
- Menge, 22, **34**
 - Leere Menge, 34
 - offene, 276
- Minimalstelle, *siehe* Minimum
- Minimierungsproblem
 - in $\mathbb{R}^2$, 178
 - in $\mathbb{R}^2$
 - Lösung, 179
- Minimum
 - im $\mathbb{R}^n$, 295
- Mittelwertsatz, 108
- Mittelwertsatz
 - der Integralrechnung, 163
- Niveaumenge, 275
- Niveaumengen, 275
- Norm
 - euklidische, 45
- Norm
 - auf reellem Vektorraum, 184
 - euklidische, 184
 - euklidische
 - in $\mathbb{R}^3$, 177
 - in $\mathbb{R}^n$, 177
 - in L_2 , 185
- Normalverteilung, 242
- normiert
 - Dichtefunktion, 238
- Notwendige Bedingung, 22, 295
- Nullfolge, 72
- Nullstellensatz, 95
- o.B.d.A., 67
- Oberintegral, 161
- Obersumme, 161
- Offene Menge, 276
- Operator
 - selbstadjungierter, 195
- Operator
 - hermitescher, 202
- Optimierung, 295
- orthogonal, 177
- orthogonal
 - Projektion
 - auf eine Gerade, 178
 - auf einen Unterraum, 182
 - in $\mathbb{R}^n$, 182
- Orthogonalbasis
 - in $\mathbb{R}^n$, 183
 - Koeffizienten, 183
- Orthogonalität, 46
- Orthogonalität
 - von Eigenvektoren, 196
- Orthogonalsystem, 182
- Orthonormalbasis, 192
- Orthonormalsystem
 - vollständiges, 189
- Oszillator
 - gedämpfter, 311
 - harmonischer, 308
- Parameterschätzung, 249
- Partielle Ableitung

- höherer Ordnung, 281
- partielle Differentialgleichung, 200
- partielle Integration, 166
- partielle Integration
 - Beispiele, 166
- Pascalsches Dreieck, 29, 30
- Permutation, 130
- Poisson-Verteilung, 234
- Polarkoordinaten, 284, 288
- positiv definit
 - Norm, 184
 - Skalarprodukt, 183
- Potential, 298
- Prävalenz, 214
- Produkt
 - von Wahrscheinlichkeitsräumen, 216
- Produktexperimente, 215
- Produktformel
 - für Wahrscheinlichkeiten, 214, 215
- Produktregel, 102
- Projektion
 - orthogonale
 - auf eine Gerade, 178
 - auf einen Unterraum, 182
 - in $\mathbb{R}^n$, 182
- Proximum
 - zu einer Geraden in $\mathbb{R}^2$, 178
- Quantentheorie, 202
- Quantenzahl, 202
- Quantil, 248
- Quelle, 301
 - Definition, 301
- Quotientenregel, 102
- radioaktiv
 - Zerfall, 234
- Rang, 53
- rationale Zahl, 71
- Realisierung, 249
- Realteil, 119
- reelle Zahl, 71
- Reihe, **81**
- absolute Konvergenz, 84
- alternierende harmonische, 84
- bedingte Konvergenz, 84
- Konvergenz, 83
- Konvergenzkriterien, 83
- Leibnizsches Kriterium, 83
- Majorante, 85
- Minorante, 85
- relative Häufigkeit, 203, 250
- Restglied
 - zum Taylorpolynom, 111
- Riemann-Integral, siehe Integral 162
- Rolle
 - Satz von, 107
- Sarrus, 135
 - Schema, 135, 135
- Satz, **25**
 - von Bolzano-Weierstraß, 78
 - von Rolle, 107
- Schätzer, 250, 251
- Schätzer
 - asymptotische Erwartungstreue, 256
 - erwartungstreu, 252
 - Erwartungswert eines Schätzers, 252
 - für Erwartungswert, 252
 - für Varianz, 253
 - Konsistenz, 256
- Schätzproblem, 250, 251
- schwaches Gesetz der großen Zahlen, 232, 233
- Sekante, 271
 - kursiv, 98
- selbstadjungierter Operator, 195
- Senke, 301
 - Definition, 301
- Sensitivität, 214
- Signum-Funktion, 131
- Singularität
 - eines Integranden, 172
- Skalarmultiplikation, 38
- Skalarprodukt, 44, 177
- Skalarprodukt
 - in reellem Vektorraum, 183

- Standard-Skalarprodukt in $\mathbb{R}^n$, 177
- Spann, 48
- Spur einer Matrix, 143
- Stabdiagramm
 - Binomialverteilung, 218
- Stabilität, 319, 322
 - asymptotische, 319, 322
 - eines Fixpunkts, 322, 323
 - eines linearen Systems, 319
- Stammfunktion, 164
- Standard-Normalverteilung, 242, 242, 247
- Standard-Normalverteilung
 - Verteilungsfunktion, 246
- Standard-Skalarprodukt
 - in $\mathbb{R}^n$, 177
- Standardabweichung
 - Definition, 224
- starkes Gesetz der großen Zahlen, 234
- Statistik, 249
- Stetigkeit, 90
 - δ - ϵ -Kriterium, 93
 - Folgenkriterium, 92
- Stichprobe, 249
- Stichprobenraum, 251
- Streuung
 - Definition, 224
- Student-Verteilung, 259
- Substitutionsregel, 167
- Substitutionsregel
 - Beispiel, 168, 169
- Sukzessive Integration, 285
- surjektiv, 36
- Symmetrie
 - Skalarprodukt, 184
- Symmetrische Gruppe, 130
- Symmetrische Matrix, 195
- t -Test, 259, 263
- t -Verteilung, 259
- Tangente
 - kursiv, 98
- Tangentialvektor, 271
- Taylorentwicklung, 111
- Exponentialfunktion, 113
- Logarithmusfunktion, 113
- Taylorpolynom, 111
 - Restglied, 111
- Taylorreihe, 112
- Test
 - diagnostischer, 212
- Theorem, 25
- theoretische Chemie, 202
- totale Ableitung, **278**
- totale Wahrscheinlichkeit
 - Formel, 209
- Totales Differentia, **279**
- Trajektorie, 304
- Treppenfunktion, 159
- Tschebyscheff-Ungleichung, 232
- Unabhängigkeit
 - von Ereignissen, 214
 - von Zufallsvariablen, 221
- unkorreliert, 225
- Unterintegral, 161
- Unterraum
 - affiner, 66
- Untersumme, 161
- Untervektorraum, **39**
- Urbild, 36
- Varianz, 222
- Varianz
 - beim Laplace-Würfel, 226
 - Definition, 224
 - der Binomialverteilung, 226
 - Eigenschaften, 225
 - einer reellen kontinuierlichen Wahrscheinlichkeitsverteilung, 239
 - einer reellen kontinuierlichen Zufallsvariablen, 238
- Vektorfeld, 296
 - konservatives, 298
 - Veranschaulichung, 297
 - zeitabhängiges, 296
 - zeitunabhängiges, 296

- Vektorprodukt, 47
- Vektorraum, **37**
 - Assoziativgesetz, 38
 - Kommutativgesetz, 38
 - Nullvektor, 43
 - Skalarmultiplikation, 38
 - Vektoraddition, 37
- Veranschaulichung
 - einer Funktion mehrerer Argumente, 275
- Verknüpfung
 - von Abbildungen, 37
- Vermutung, 25
- Verteilung
 - einer Zufallsvariablen, 221
 - gemeinsame, 226
- Verteilungsfunktion
 - zu einer Wahrscheinlichkeitsdichte, 238
- Vertrauensintervall, 257
- vollständig
 - Orthonormalsystem, 189
- Wärmeleitungsgleichung, 200
- Würfel
 - fairer, 203
- Wahrheitstafel, 20
- Wahrscheinlichkeit, 203, 204
- Wahrscheinlichkeit
 - bedingte, 207
- Wahrscheinlichkeitsbaum, 211
- Wahrscheinlichkeitsdichte, 237
- Wahrscheinlichkeitsfunktion, 205
- Wahrscheinlichkeitsmaß, 204
- Wahrscheinlichkeitsmaß
 - zu einer Wahrscheinlichkeitsdichte, 238
- Wahrscheinlichkeitsraum
 - endlicher, 204
 - kontinuierlicher, 237
 - Laplacescher, 205
 - unendlicher, 234
 - unendlicher diskreter, 234
- Wahrscheinlichkeitstheorie, 203
- Wahrscheinlichkeitsverteilung, 204
- Wartezeit
 - beim Poisson-Prozeß, 241
- Wasserströmung, 296
- Wegintegral, *siehe* Kurvenintegral
- Welle
 - longitudinal, 302
 - transversal, 302
- Wellenfunktion, 202
- Wertemenge, 91
- Wetterkarte, 302
- Widerspruchsbeweis, 26
- Windkanal, 297
- Winkel
 - zwischen Vektoren, 181
- Zahl
 - komplex, 119
 - rational, 71
 - reell, 71
- zentraler Grenzwertsatz, 245
- Zerfall
 - radioaktiv, 234
- Zerlegung
 - eines Intervalls, 161
- Zufallsexperiment, 203
- Zufallsvariable, 219
- Zufallsvariable
 - reelle, 219
 - reelle kontinuierliche, 238
- Zustandsvektor, 304
- Zwischenwertsatz, 95