
Mathematik für Informatiker I

Wintersemester 2003, Prof. J. Weickert

erstellt von: Rico Philipp, Kai Hagenburg

Version vom: 21. Juli 2004

Inhaltsverzeichnis

I	Diskrete Mathematik	22
1	Mengen	23
1.1	Menge	23
1.2	Anmerkung:	23
1.3	Teilmenge	24
1.4	Potenzmenge	24
1.5	Operationen mit Mengen	25
1.6	Rechenregeln für Durchschnitte und Vereinigungen	26
1.7	Rechenregeln für Komplementbildung	27
1.8	Unendliche Durchschnitte und Vereinigungen	27
1.9	Beispiel:	28
1.10	Kartesisches Produkt	28
1.11	Beispiel:	28
1.12	Verallgemeinerung	28
2	Aussagenlogik	29
2.1	Bedeutung in der Informatik	29
2.2	Aussage	29
2.3	Beispiele	29
2.4	Verknüpfungen von Aussagen	30
2.5	Anmerkung:	30
2.6	Tautologie	30

2.7	Beispiele	31
2.8	wichtige Tautologien	31
2.9	Beispiele	31
2.10	Quantoren	32
2.11	Negation von Aussagen	33
3	Beweisprinzipien	35
3.1	Bedeutung in der Informatik	35
3.2	Direkter Beweis	35
3.3	Beispiel:	36
3.4	Beweis durch Kontraposition	36
3.5	Widerspruchsbeweise	36
3.6	Beispiel:	36
3.7	Äquivalenzbeweise	37
3.8	Beweis durch vollständig Induktion	37
3.9	Beispiel:	37
3.10	Anmerkung:	38
4	Relationen	39
4.1	Motivation	39
4.2	Relation	39
4.3	Beispiel:	40
4.4	Reflexivität, Symmetrie, Transitivität bei Relation	40
4.5	Beispiele	40
4.6	Äquivalenzklassen	41
4.7	Beispiel:	41
4.8	Partition	42
4.9	Partitionseigenschaften von Äquivalenzklassen	42
4.10	Teilrelation	43
4.11	Beispiele	43

5	Abbildungen	45
5.1	Abbildung	45
5.2	Beispiel:	45
5.3	Wichtige Begriffe	45
5.4	Definitionsbereich	46
5.5	Beispiel:	46
5.6	surjektiv, injektiv, bijektiv	47
5.7	Beispiele	47
5.8	Umkehrabbildung	48
5.9	Beispiele	48
5.10	Komposition	48
5.11	Assoziativität der Verknüpfung	48
5.12	<u>Vorsicht</u>	49
5.13	Kriterium für Injektivität, Surjektivität, Bijektivität	49
5.14	Mächtigkeit von Mengen	50
5.15	Beispiele	50
5.16	Abzählbarkeit von $\mathbb{Q}$	50
5.17	Äquivalenz von Surjektivität und Injektivität	51
5.18	Folgerungen	52
5.19	Schubfachprinzip	52
5.20	Beispiele	53
6	Primzahlen und Teiler	55
6.1	Bedeutung in der Informatik	55
6.2	Division mit Rest	55
6.3	Teilbarkeit, Teiler, Primzahl	56
6.4	Beispiele	56
6.5	Teilbarkeitsregeln	56
6.6	Fundamentalsatz der Zahlentheorie	56
6.7	Sieb des Erathostenes (Primzahlfaktorisationen)	57

6.8	Beispiel:	57
6.9	größter gemeinsamer Teiler (ggT)	57
6.10	Eigenschaften des ggT	58
6.11	Euklidischer Algorithmus	59
6.12	Beispiel:	59
6.13		59
7	Modulare Arithmetik	61
7.1	Bedeutung in der Informatik	61
7.2	kongruent modulo m, Modul	61
7.3	Zusammenhang Kongruenz - Division mit Rest	61
7.4	Interpretation als Äquivalenzrelation	62
7.5	Addition von Elementen zweier Restklassen	63
7.6	modulare Addition	63
7.7	Beispiel: Additionstafeln in $\mathbb{Z}_6$	63
7.8	Eigenschaften der modularen Addition	64
7.9	Multiplikation von Elementen zweier Restklassen	65
7.10	modulare Multiplikation	65
7.11	Beispiel: Multiplikationstafeln in $\mathbb{Z}_6$	65
7.12	Eigenschaften der modularen Multiplikation	66
7.13	Multiplikative inverse Elemente in $\mathbb{Z}_m$	66
II	Algebra	68
8	Gruppen	69
8.1	Bedeutung in der Informatik	69
8.2	Gruppe	69
8.3	Beispiele	70
8.4	Eindeutigkeit des neutr. Elements und der inversen Elemente	71
8.5	Untergruppenkriterium	72

8.6	Beispiel:	72
8.7	Zyklenschreibweise für Permutationen	73
8.8	Nebenklassen, Rechts- und Linksnebenklassen	73
8.9	Nebenklassenzerlegung einer Gruppe	74
8.10	Beispiele	74
8.11	G modulo U	75
8.12	von Lagrange	75
8.13	Beispiel:	75
8.14		76
8.15	Beispiel:	76
8.16	Abbildung zwischen Gruppen	76
8.17	Bild, Kern	77
8.18	Homomorphismus für Gruppen	77
9	Ringe	79
9.1	Motivation	79
9.2	Ring	79
9.3	Konventionen	80
9.4	Beispiele	80
9.5	Unterringkriterium	80
9.6	Beispiel:	81
9.7	Polynomringe	81
9.8	Horner-Schema	82
9.9	Polynomdivision	82
9.10	Praktische Durchführung der Polynomdivision	83
9.11	Abspaltung von Nullstellen	83
9.12	Zahl der Nullstellen	84
9.13	ggT von Polynomen	84
9.14	Euklidischer Algorithmus für Polynome	84
9.15	reduzibel, irreduzibel	85

9.16 Beispiele	85
9.17 Irreduzible Polynome als „Primfaktoren“ in $(\mathbb{R}[x], +, \cdot)$	85
9.18 Beispiel:	85
10 Körper	87
10.1 Körper	87
10.2 Bemerkung:	87
10.3 Beispiele	87
10.4 Eigenschaften von Körpern	88
10.5 Endliche Körper	88
10.6 Der Körper der komplexen Zahlen	89
10.7 Komplexe Zahlen	89
10.8 Konsequenzen aus Def. 10.7	90
10.9 Körper der komplexen Zahlen	90
10.10 Praktisches Rechnen mit komplexen Zahlen	90
10.11 komplex konjugiertes Element, Betrag, Realteil, Imaginärteil	91
10.12 Geometrische Interpretation	91
10.13 Nützlichkeit des konjugiert komplexen Elementes	92
10.14 Fundamentalsatz der Algebra	92
10.15 Beispiel:	92
10.16 Wozu sind komplexe Zahlen noch nützlich?	93
11 Polynome auf allgemeinen Körpern	95
11.1 Übertragung von Resultaten aus $\mathbb{R}[x]$	95
11.2 Beispiele	95
11.3 Anwendung der Polynomdivision in $\mathbb{Z}_2[x]$ zur Datensicherung	96
12 Boole'sche Algebren	99
12.1 Motivation	99
12.2 Boole'sche Algebra	99
12.3 Beispiel 1: Operationen auf Mengen	100

12.4 Beispiel 2: Aussagenlogik	100
12.5 Eigenschaften Boole'scher Algebren	101
12.6 Endliche Boole'sche Algebren	101
13 Axiomatik der reellen Zahlen	103
13.1 Motivation	103
13.2 angeordnet, Positivbereich	103
13.3 Ordnungsbegriffe	104
13.4 Eigenschaften angeordneter Körper	104
13.5 Konsequenzen	105
13.6 Maximum und Minimum	105
13.7 Beispiel:	105
13.8 Infimum, Supremum	106
13.9 Beispiel:	106
13.10 vollständig	107
13.11 Axiomatische Charakterisierung von $\mathbb{R}$	107
13.12 Eigenschaften der reellen Zahlen	107
13.13 Betragsfunktion	107
13.14 Eigenschaften des Betrags	107
III Eindimensionale Analysis	109
14 Folgen	111
14.1 Motivation	111
14.2 Folge	111
14.3 Beispiel:	111
14.4 (streng) monoton fallend / wachsend; beschränkt	112
14.5 ε - Umgebung	112
14.6 konvergente Folge	112
14.7 Beispiel:	113

14.8	uneigentliche Konvergenz, bestimmte Divergenz	113
14.9	Beschränktheit konvergenter Folgen	113
14.10	Eindeutigkeit des Grenzwerts	114
14.11	Konvergenzkriterien	114
14.12	Beispiel:	115
14.13	Rechenregeln für Grenzwerte	115
14.14	Beispiel:	115
14.15	Der Grenzwert $\lim_{n \rightarrow \infty} \left(1 + \frac{p}{n}\right)^n$	116
15	Landau - Symbole	119
15.1	Motivation	119
15.2	Groß-O, Klein-o	119
15.3	Beispiel:	119
15.4	Mehrdeutigkeit	120
15.5	Rechenregeln für Landau-Symbole	120
15.6	Vergleichbarkeit von Folgen	121
15.7	$O(A) = O(B), O(A) < O(B)$	121
15.8	Häufig verwendete Prototypen von Vergleichsfunktionen	121
15.9	Vergleich von Ordnungen	122
16	Reihen	123
16.1	Motivation	123
16.2	Reihe, n-te Partialsumme	123
16.3	Konvergenzkriterien für Reihen	124
16.4	Beispiel:	125
16.5	absolut konvergente Reihe	126
16.6	Beispiel:	127
16.7	Kriterien für absolute Konvergenz	127
16.8	Bemerkung:	128
16.9	Beispiel:	129

16.10 Umordnungssatz	130
16.11 Bemerkung:	131
16.12 Umordnung nicht absolut konvergenter Reihen	131
16.13 Produkt von Reihen	132
16.14 Folgerung:	133
16.15 Beispiel:	133
17 Potenzreihen	135
17.1 Motivation	135
17.2 Potenzreihe, Entwicklungspunkt	135
17.3 Beispiel:	135
17.4 Für eine Potenzreihe, die nicht für alle $x \in \mathbb{R}$ konvergiert, gibt es ein $R \geq 0$, so dass sie für $ x < R$ absolut konvergiert und für $ x > R$ divergiert.	136
17.5 Anmerkung:	136
17.6 Hadamard, Berechnung des Konvergenzradius	136
17.7 Beispiel:	136
18 Darstellung von Zahlen in Zahlssystemen	139
18.1 Motivation	139
18.2 Darstellung natürlicher Zahlen in Zahlssystemen	139
18.3 Ziffern, Basis	139
18.4 Beispiel:	140
18.5 b-adischer Bruch	140
18.6 Beispiel:	141
18.7 Darstellung reeller Zahlen in Zahlensystemen	141
18.8 Bemerkung:	142
19 Binomialkoeffizienten und die Binomialreihe	143
19.1 Motivation	143
19.2 Binomialkoeffizient	143
19.3 Rekursionsbeziehung für Binomialkoeffizienten	144

19.4 Pascal'sches Dreieck	144
19.5 Direkte Formel für Binomialkoeffizienten	144
19.6 Beispiel:	145
19.7 Binomialsatz	145
19.8 Beispiel:	145
19.9 Beweis des Binomialsatzes	145
19.10 Binomialreihe: Motivation	146
19.11 Binomialreihe	147
19.12 Konvergenz der Binomialreihe	147
19.13 Beispiele	148
20 Stetigkeit	149
20.1 Motivation	149
20.2 Konvergenz gegen Grenzwert	149
20.3 Beispiele	149
20.4 Grenzwertsätze für Funktionen	150
20.5 Stetigkeit	150
20.6 $\epsilon - \delta$ -Kriterium der Stetigkeit	150
20.7 Bemerkung:	151
20.8 Beispiel:	152
20.9 Eigenschaften stetiger Funktionen	152
20.10 Gleichmäßige Stetigkeit	154
21 Wichtige stetige Funktionen	157
21.1 Motivation	157
21.2 Potenzfunktionen	157
21.3 Wurzelfunktionen	158
21.4 Exponentialfunktion	158
21.5 Logarithmusfunktionen	160
21.6 Trigonometrische Funktionen	160

21.7	Trigonometrische Umkehrfunktionen	163
22	Differenzierbarkeit	165
22.1	Motivation	165
22.2	differenzierbar, Ableitung, Differentialquotient	166
22.3	Bemerkungen:	166
22.4	Beispiel:	167
22.5	Ableitungen elementarer Funktionen	168
22.6	Differentiationsregeln	168
22.7	Beispiele	169
22.8	Bemerkung:	170
22.9	Zweite Ableitung	170
22.10	Beispiel:	171
23	Mittelwertsätze und die Regel von L'Hospital	173
23.1	Motivation	173
23.2	Mittelwertsätze	173
23.3	Regel von l'Hospital	175
23.4	Beispiel:	176
24	Der Satz von Taylor	177
24.1	Motivation	177
24.2	Satz von Taylor	177
24.3	Bemerkung:	179
24.4	Beispiel:	179
24.5	Taylorreihen	180
24.6	Bemerkungen:	180
24.7	Beispiele	180
24.8	Gliedweise Differentiation von Potenzreihen	181
24.9	Beispiel:	181

25 Geometrische Bedeutung der Ableitungen	183
25.1 Motivation	183
25.2 Monotonie	183
25.3 Monotonie differenzierbarer Funktionen	183
25.4 Beispiele	184
25.5 Konvexität, Konkavität	185
25.6 Veranschaulichung:	185
25.7 Konvexität / Konkavität von C^1 -Funktionen	185
25.8 Konvexität / Konkavität von C^2 -Funktionen	186
25.9 Beispiele	186
25.10 Lokale Minima und Maxima	186
25.11 Notwendige Bedingung für lokale Extrema	187
25.12 Beispiele	187
25.13 Hinreichende Bedingung für lokale Extrema	188
25.14 Beispiele	189
25.15 Wendepunkt	189
25.16 Bemerkung:	190
25.17 Notwendige Bedingung für Wendepunkte	190
25.18 Hinreichende Bedingung für Wendepunkte	190
25.19 Beispiel:	191
25.20 Kurvendiskussion	191
 26 Fixpunkt-Iteration	 195
26.1 Motivation	195
26.2 Beispiel:	196
26.3 Lipschitz-stetig, Lipschitz-Konstante	197
26.4 Bemerkungen:	197
26.5 Lipschitz-Stetigkeit von C^1 -Funktionen	197
26.6 Banachscher Fixpunktsatz, 1D-Version	198
26.7 Bemerkung:	199

26.8 Beispiel:	199
27 Das Newton-Verfahren	201
27.1 Motivation	201
27.2 Newton-Verfahren für eine nichtlineare Gleichung	201
27.3 Beispiele	202
27.4 Bemerkung:	202
27.5 Konvergenz des Newtonverfahrens	203
27.6 Bemerkung:	204
28 Das Bestimmte Integral	205
28.1 Motivation	205
28.2 Zerlegung, Partition, Feinheit	205
28.3 Riemann-Summe	206
28.4 Bemerkung:	206
28.5 Riemann-Integral	207
28.6 Beispiele	207
28.7 Eigenschaften des Riemann-Integrals	208
28.8 Bemerkung:	208
28.9 Monotonie des Riemann-Integrals	209
28.10Folgerung:	209
28.11Folgerung (Nichtnegativität):	209
28.12Linearität der Integration	210
28.13Zusammensetzung von Integrationsintervallen	211
28.14Riemann'sches Kriterium	211
28.15Integrierbarkeit monotoner Funktionen	212
28.16Integrierbarkeit stetiger Funktionen	212
28.17Beispiele	213
28.18Mittelwertsatz der Integralrechnung	213
28.19Korollar: (<i>Integralmittel</i>)	214

28.20	Veranschaulichung	214
28.21	Volumen von Rotationskörpern	215
28.22	Beispiel:	215
29	Das unbestimmte Integral und die Stammfunktion	217
29.1	Motivation	217
29.2	Stammfunktion	217
29.3	Hauptsatz der Differential- und Integralrechnung	217
29.4	Unbestimmte Integral	218
29.5	Beispiele von unbestimmten Integralen	219
29.6	Integrationsregeln	219
29.7	Beispiele zur partiellen Integration	220
29.8	Beispiele zur Substitutionsregel	221
29.9	Bemerkung:	222
29.10	Integration rationaler Funktionen	222
30	Uneigentliche Integrale	223
30.1	Motivation	223
30.2	Unendliche Integrationsgrenzen	223
30.3	Beispiele	224
30.4	Konvergenz	224
30.5	Beispiel:	225
30.6	Bemerkung:	225
30.7	Konvergenzkriterien für uneigentliche Integrale	225
30.8	Beispiel: <u>Das Dirichlet-Integral $\int_0^\infty \frac{\sin x}{x} dx$</u>	226
30.9	Beispiel: <u>Die Gammafunktion (Euler, 1707-1783)</u>	226
31	Numerische Verfahren zur Integration	229
31.1	Motivation	229
31.2	Einfachstes Beispiel: Rechteckregel	230

31.3	Zweiteinfaches Beispiel: Trapezregel	230
31.4	Bemerkung:	230
31.5	Euler-Maclaurinsche Summenformel	231
31.6	Romberg-Quadratur	232
31.7	Beispiel:	232
32	Kurven und Bogenlänge	235
32.1	Motivation	235
32.2	Parametrisierte Kurve	235
32.3	Beispiele	236
32.4	Parameterwechsel, Umparametrisierung	237
32.5	Bemerkung:	237
32.6	Länge, Rektifizierbarkeit einer Kurve	238
32.7		239
32.8	Beispiel:	240
32.9	Parametrisierungsinvarianz	240
32.10	Bogenlängenfunktion	241
32.11	Tangentialvektor, Tangenteneinheitsvektor	241
32.12	Beispiel:	242
32.13	Die Sektorfläche	243
32.14	Orientierter Flächeninhalt	244
32.15	Sektorformel von LEIBNIZ	244
32.16	Beispiele	245
33	Mehrfachintegrale	247
33.1	Motivation	247
33.2	Zerlegung, Feinheit, Flächeninhalt	247
33.3	Bemerkungen:	248
33.4	Unterintegral, Oberintegral, Riemann-Integral	249
33.5	Eigenschaften des Mehrfachintegrals	249

33.6 Satz von Fubini	250
33.7 Beispiel:	251
33.8 Bemerkungen:	251
 IV Lineare Algebra	 252
 34 Vektorräume	 253
34.1 Motivation	253
34.2 K -Vektorraum	253
34.3 Rechenregeln für Vektorräume	254
34.4 Beispiele	254
34.5 Unterraum, Untervektorraum, Teilraum	255
34.6 Unterraumkriterium	255
34.7 Beispiele	256
34.8 Linearkombination, Erzeugnis	256
34.9 Erzeugnis als Unterraum	256
34.10 Beispiel:	257
34.11 Lineare Abhängigkeit	257
34.12 Kriterium für lineare Unabhängigkeit	258
34.13 Lineare Unabhängigkeit	258
34.14 Beispiel:	258
34.15 Basis	259
34.16 Beispiele	259
34.17 Existenz einer Basis	260
34.18 Austauschsatz	260
34.19 Eindeutigkeit der Zahl der Basisvektoren	261
34.20 Dimension	262
34.21 Basiskriterium linear unabhängiger Mengen	262
 35 Lineare Abbildungen	 263

35.1 Motivation	263
35.2 lineare Abbildung, Vektorraumhomomorphismus	263
35.3 Beispiele	264
35.4 Bild, Kern, Monomorphismus,...	264
35.5 Eigenschaften linearer Abbildungen	264
35.6 Beispiel:	265
35.7 Eindeutigkeit der Darstellung in einer festen Basis	265
35.8 Umrechnung von Koordinatendarstellungen	266
35.9 Charakterisierung einer linearen Abbildung durch ihre Wirkung auf die Basis	267
35.10 Isomorphie endlichdimensionaler Vektorräume	267
36 Matrixschreibweise für lineare Abbildungen	269
36.1 Motivation	269
36.2 Struktur linearer Abbildungen zwischen endlichdimensionalen Vek- torräumen	269
36.3 $(m \times n)$ -Matrix, Zeilenindex, Spaltenindex	270
36.4 Matrix - Vektor - Produkt	271
36.5 Beispiele für lineare Abbildung in Matrixschreibweise	271
36.6 Skalare Multiplikation, Matrixoperationen	273
36.7 Beispiele	273
36.8 Eigenschaften der Matrixoperationen	273
36.9 Bemerkung:	274
36.10 Inverse Matrix	275
36.11 Invertierbarkeit einer Matrix	275
36.12 Gruppeneigenschaft von $GL(n, K)$	275
36.13 Bemerkungen:	276
37 Der Rang einer Matrix	277
37.1 Motivation	277
37.2 Spaltenrang	277

37.3 Aussagen über den Rang einer Matrix	277
37.4 Beispiele	278
37.5 Transponierte Matrix	278
37.6 Beispiel:	278
37.7 Gleichheit von Spaltenrang und Zeilenrang	279
38 Gauss'scher Algorithmus und lineare Gleichungssysteme	281
38.1 Motivation	281
38.2 Idee	282
38.3 elementare Zeilenumformungen	282
38.4 Bei elementaren Zeilenumformungen bleibt der Rang einer Matrix erhalten.	282
38.5 Gauss'scher Algorithmus	283
38.6 Beispiel:	285
38.7 Der Rang einer Matrix in Zeilen-Stufen-Form ist gleich der An- zahl ihrer Leitkoeffizienten.	285
38.8 Satz:	286
38.9 Umformung einer invertierbaren Matrix zur Einheitsmatrix . . .	287
38.10 Berechnung der inversen Matrix	287
38.11 Beispiel	288
38.12 Lineare Gleichungssysteme	289
38.13 Lösungsverhalten linearer Gleichungssysteme	289
38.14 Bemerkungen:	290
38.15 Invarianz der Lösungsmenge unter Matrixmultiplikation	290
38.16 Invarianz der Lösungsmenge unter elementaren Zeilenumformungen	291
38.17 Gauß-Algorithmus zum Lösen linearer Gleichungssysteme	291
38.18 Geometrische Interpretation linearer Gleichungssysteme	293
39 Iterative Verfahren für lineare Gleichungssysteme	295
39.1 Motivation	295
39.2 Grundstruktur klassischer iterativer Verfahren	296

39.3 Das Jacobi-Verfahren (Gesamtschrittverfahren)	296
39.4 Beispiel	297
39.5 Das Gauß-Seidel-Verfahren (Einzelschrittverfahren)	297
39.6 Beispiel:	298
39.7 Das SOR-Verfahren (SOR = successive overrelaxation)	298
39.8 Konvergenzresultate	299
39.9 Effizienz	299
40 Determinanten	301
40.1 Motivation	301
40.2 Determinantenfunktion, Determinante	301
40.3 Eindeutigkeit der Determinante	302
40.4 Determinanten einer 2×2 -Matrix	302
40.5 Unterdeterminante	303
40.6 Beispiel	303
40.7 Laplace'scher Entwicklungssatz	303
40.8 Beispiel:	304
40.9 Rechenregel für Determinanten	304
40.10 Bedeutung der Determinanten	305

Teil I

Diskrete Mathematik

Kapitel 1

Mengen

Der Mengenbegriff ist von grundlegender Bedeutung in vielen Gebieten der Informatik, z.B. Datenbanken.

1.1 Definition: (*Menge*)

Unter einer Menge M verstehen wir jede Zusammenfassung von bestimmten wohlunterscheidbaren Objekten der Anschauung oder des Denkens, welche die Elemente von M genannt werden, zu einem Ganzen.

1.2 Anmerkung:

- (a) Dieser Mengenbegriff geht auf Georg Cantor zurück (1845-1918). Er begründete die moderne Mengenlehre.
- (b) Er ist nicht unumstritten und auch nicht widerspruchsfrei, genügt jedoch ausreichend für unsere Anwendungen
Beispiel: Der Barbier rasiert alle Menschen, die sich nicht selbst rasieren können. Darf er sich selbst rasieren?
- (c) Die Elemente einer Menge werden in geschweiften Klammern eingeschlossen.
Beispiel: $M = \{4, 1, 5\}$
- (d) Die Reihenfolge der Elemente spielt keine Rolle, und Elemente dürfen mehrfach auftreten. Wir unterscheiden also nicht zwischen $\{4, 1, 5\}$ und $\{1, 4, 5\}$.
- (e) Symbole für wichtige Mengen:
 - (i) $\mathbb{N} :=^1 \{x \mid^2 x \text{ ist eine natürliche Zahl}\} = \{1, 2, 3, \dots\}$

¹„ist definiert als“

²„mit der Eigenschaft“

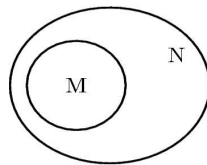
- (ii) $\mathbb{N}_0 := \{0, 1, 2, 3, \dots\}$ natürliche Zahlen mit Null
- (iii) $\mathbb{Z} := \{0, 1, -1, 2, -2, \dots\}$ ganze Zahlen
- (iv) $\mathbb{Q} := \{x \mid x = \frac{p}{q}, p \in \mathbb{Z}, q \in \mathbb{N}\}$ rationale Zahlen
- (v) $\mathbb{R}$ reelle Zahlen
- (vi) $\emptyset, \{\}$ = leere Menge (enthält kein Element)
- (f) Mengen können wieder Mengen enthalten.
Beispiel: $M = \{\mathbb{N}, \mathbb{Z}, \mathbb{Q}\}$ **dreielementige** Elemente.

1.3 Definition: (Teilmenge)

M heißt Teilmenge von N ($M \subset N, N \supset M$), wenn jedes Element in M auch Element in N ist:

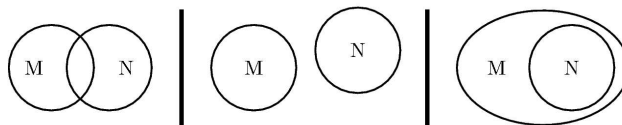
$$x \in M \Rightarrow x \in N.$$

Visualisierung durch sogenannte Venn-Diagramme: Ist M nicht Teilmenge von



N , schreibt man $M \not\subset N$.

$\emptyset$ ist Teilmenge jeder Menge M . 2 Mengen N, M sind gleich ($M = N$), falls



$N \subset M$ und $M \subset N$ (wichtig für Beweise). Falls $N \subset M$ und $N \neq M$, dann schreibt man auch $N \subsetneq M$ („echt enthalten“).

1.4 Definition: (Potenzmenge)

Ist M eine Menge, so heißt $P(M) := \{X \mid X \subset M\}$ Potenzmenge von M .

Beispiel:

$$(a) \quad M = \{1, 2\} \Rightarrow P(M) = \{\emptyset, \{1\}, \{2\}, \{1, 2\}\}$$

- (b) $M = \{a, b, c\} \Rightarrow P(M) = \{\emptyset, \{a\}, \{b\}, \{c\}, \{a, b\}, \{a, c\}, \{b, c\}, \{a, b, c\}\}$
(c) $M = \emptyset \Rightarrow P(M) = \{\emptyset\}$

Die Potenzmenge einer n-elementigen Menge enthält 2^n **Elemente**.

1.5 Operationen mit Mengen

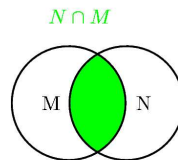
- (a) **Durchschnitt** zweier Mengen M und N :

$$M \cap N := \{x \mid x \in M \text{ und } x \in N\}$$

N und M heißen disjunkt, falls $N \cap M \neq \emptyset$.

Beispiel: $M = \{1, 3, 5\}$, $N = \{2, 3, 5\}$, $S = \{5, 7, 8\}$.

$$M \cap N = \{3, 5\} \quad (M \cap N) \cap S = \{5\}$$

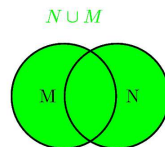


- (b) **Vereinigung** zweier Mengen

$$M \cup N := \{x \mid x \in M \text{ oder } x \in N\}$$

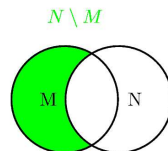
Dabei darf x auch in beiden Mengen sein.

(„oder“ ist kein „exklusives oder“)



- (c) **Differenz** zweier Mengen

$$M \setminus N := M - N := \{x \mid x \in M \text{ und } x \notin N\}$$



$M \setminus N$ heißt Komplement von N in M . Schreibweise: $\overline{N}^M$.
Wenn Grundmenge M klar ist, schreibt man auch: $\overline{N}$, $\complement N$, N^C .

1.6 Satz: (Rechenregeln für Durchschnitte und Vereinigungen)

Seien M , N und S Mengen. Dann gelten folgende Gesetze:

(a) Kommutativgesetz

$$(i) M \cup N = N \cup M$$

$$(ii) M \cap N = N \cap M$$

(b) Assoziativgesetz

$$(i) (M \cup N) \cup S = M \cup (N \cup S)$$

$$(ii) (M \cap N) \cap S = M \cap (N \cap S)$$

(c) Distributivgesetz

$$(i) M \cap (N \cup S) = (M \cap N) \cup (M \cap S)$$

$$(ii) M \cup (N \cap S) = (M \cup N) \cap (M \cup S)$$

Ferner gilt für jede Menge M :

- $M \cup \emptyset = M$
- $M \cap \emptyset = \emptyset$
- $M \setminus \emptyset = M$

Beweis:

von (c)(i)

$$\text{Zeige } \underbrace{M \cap (N \cup S)}_{=:A} = \underbrace{(M \cap S) \cup (M \cap N)}_{=:B}$$

Um $A = B$ zu zeigen, zeigen wir $A \subseteq B \wedge A \supseteq B$

„ $A \subset B$ “: Sei $x \in A$. $\Rightarrow x \in M$ und $(x \in N \text{ oder } x \in S)$.

- Falls $x \in N : x \in M \cap N$
 $\Rightarrow x \in (M \cap N) \cup (M \cap S) = B$
- Falls $x \in S : x \in M \cap S$
 $\Rightarrow x \in (M \cap N) \cup (M \cap S) = B$

„ $B \subset A$ “: Sei $x \in B$.

$\Rightarrow x \in M \cap N$ oder $x \in M \cap S$

- Falls $x \in M \cap N$:
 $\Rightarrow x \in M$ und $x \in N$
 $\Rightarrow x \in N \cup S$ (wegen $x \in N$)
 $\Rightarrow x \in M \cap (N \cup S) = A$ (wegen $x \in M$).

- Falls $x \in M \cap S$:
 $\Rightarrow x \in M$ und $x \in S$
 $\Rightarrow x \in N \cup S$ (wegen $x \in S$)
 $\Rightarrow x \in M \cap (N \cup S) = A$ (wegen $x \in M$).

□³

1.7 Satz: (Rechenregeln für Komplementbildung)

Bei Rechnungen mit einer Grundmenge G gilt:

- (a) $M \setminus N = M \cap \overline{N}$
- (b) de Morgan'sche Regeln:
 - (i) $\overline{M \cup N} = \overline{M} \cap \overline{N}$
 - (ii) $\overline{M \cap N} = \overline{M} \cup \overline{N}$
- (c) Aus $M \subset N$ folgt $\overline{N} \subset \overline{M}$

(b) und (c): „Komplementbildung kehrt die Operationen um“.
Beweis von (b) $\Rightarrow$ Übungsaufgabe.

1.8 Definition: (Unendliche Durchschnitte und Vereinigungen)

Sei M eine Menge von Indizes (z.B. $M = \mathbb{N}$). Für alle $k \in M$ sei die Menge A_k gegeben. Dann gilt:

$$\begin{aligned}\bigcup_{k \in M} A_k &:= \{x \mid x \in A_i \text{ für (mind.) ein } i \in M\} \\ \bigcap_{k \in M} A_k &:= \{x \mid x \in A_i \text{ für alle } i \in M\}.\end{aligned}$$

Für endliche Mengen, z.B. $M = \{1, \dots, n\}$ stimmt diese Definition mit der bisherigen überein. Man schreibt auch

$$\begin{aligned}\bigcup_{k=1}^n A_k &:= A_1 \cup A_2 \cup \dots \cup A_n \\ \bigcap_{k=1}^n A_k &:= A_1 \cap A_2 \cap \dots \cap A_n.\end{aligned}$$

Wegen der Assoziativität ist keine Klammerung nötig.

³Beweisende

1.9 Beispiel:

$$M := \mathbb{N}, A_k := \{x \in \mathbb{R} \mid 0 < x < \frac{1}{k}\}$$

$$\bigcap_{k=1}^n A_k = A_n \quad \text{endliches Intervall}$$

$$\text{Aber: } \bigcap_{k \in \mathbb{N}} A_k = \emptyset.$$

(Es gibt keine positive Zahl x , die für alle $k \in \mathbb{N}$ in $(0, \frac{1}{k})$ liegt).

1.10 Definition: (Kartesisches Produkt)

Seien M und N Mengen, dann heißt die Menge $M \times N := \{(x, y) \mid x \in M, y \in N\}$ aller geordneten(!) Paare (x, y) mit $x \in M, y \in N$ das kartesische Produkt von M und N .

Beachte: $(3, 5) \neq (5, 3)$ (Im Unterschied zu $\{5, 3\} = \{3, 5\}$)

1.11 Beispiel:

$$M := \{1, 2\} \quad N := \{a, b, c\}$$

$$M \times N \quad ; = \quad \{(1, a), (1, b), (1, c), (2, a), (2, b), (2, c)\}$$

$$N \times M \quad ; = \quad \{(a, 1), (a, 2), (b, 1), (b, 2), (c, 1), (c, 2)\} \neq M \times N$$

Im Allgemeinen ist also $M \times N \neq N \times M$.

1.12 Verallgemeinerung

Das n-fache kartesische Produkt der Mengen $M_1, M_2, \dots, M_n$ ist gegeben durch

$$M_1 \times M_2 \times \dots \times M_n := \{(m_1, \dots, m_n) \mid m_1 \in M_1, \dots, m_n \in M_n\}.$$

Sind alle Mengen gleich, so kürzt man ab:

$$M^n := \underbrace{M \times M \times \dots \times M}_{n\text{-mal}}.$$

Beispiel: $\mathbb{R}^3 = \mathbb{R} \times \mathbb{R} \times \mathbb{R}$.

Kapitel 2

Aussagenlogik

*„Mein teurer Freund, ich rat Euch drum,
Zuerst Collegium Logicum.“*

Goethe, Faust I

2.1 Bedeutung in der Informatik

- Schaltkreisentwurf
- Verifikation
- automatisches Beweisen
- Anfragen an Suchmaschinen im Internet

2.2 Definition: *(Aussage)*

Eine Aussage ist ein Satz, in einer menschlichen oder künstlichen Sprache, dem eindeutig einer der Wahrheitswerte wahr (1) oder falsch (0) zugeordnet werden kann.

2.3 Beispiele

- (a) „Saarbrücken liegt am Rhein“ (falsch)
- (b) „ $3 < 5$ “ (wahr)

2.4 Verknüpfungen von Aussagen

Wir definieren die Verknüpfungen

- (a) $\neg \mathcal{A}$: „nicht $\mathcal{A}$ “ (**Negation**)
- (b) $\mathcal{A} \wedge \mathcal{B}$: „ $\mathcal{A}$ und $\mathcal{B}$ “ (**Konjunktion**)
- (c) $\mathcal{A} \vee \mathcal{B}$: „ $\mathcal{A}$ oder $\mathcal{B}$ “ (**Disjunktion**)
- (d) $\mathcal{A} \Rightarrow \mathcal{B}$: „ $\mathcal{A}$ impliziert $\mathcal{B}$ “ (**Implikation**)
- (e) $\mathcal{A} \Leftrightarrow \mathcal{B}$: „ $\mathcal{A}$ ist äquivalent $\mathcal{B}$ “ (**Äquivalenz**)

mittels Wahrheitstafel

$\mathcal{A}$	$\mathcal{B}$	$\neg \mathcal{A}$	$\mathcal{A} \wedge \mathcal{B}$	$\mathcal{A} \vee \mathcal{B}$	$\mathcal{A} \Rightarrow \mathcal{B}$	$\mathcal{A} \Leftrightarrow \mathcal{B}$
1	1	0	1	1	1	1
1	0	0	0	1	0	0
0	1	1	0	1	1	0
0	0	1	0	0	1	1

2.5 Anmerkung:

- (a) $\mathcal{A} \vee \mathcal{B}$ auch wahr, wenn $\mathcal{A}$ und $\mathcal{B}$ wahr sind (kein „entweder oder“).
- (b) Die Implikation $\mathcal{A} \Rightarrow \mathcal{B}$ ist immer wahr, wenn die Prämisse (Aussage $\mathcal{A}$) falsch ist.
Aus falschen Aussagen können also auch wahre Aussagen folgen:
Beispiel: „1 = -1“ falsch. Quadrieren liefert die wahre Aussage „1 = 1“.
- (c) $\mathcal{A} \Rightarrow \mathcal{B}$ ist nicht das selbe wie $\mathcal{B} \Rightarrow \mathcal{A}$.
Beispiel: $\mathcal{A}$: „Das Tier ist ein Dackel“, $\mathcal{B}$: „Das Tier ist ein Hund“.
 $\mathcal{A}$ impliziert $\mathcal{B}$, aber $\mathcal{B}$ impliziert nicht notwendigerweise $\mathcal{A}$.

2.6 Definition: (*Tautologie*)

Ein logischer Ausdruck ist eine Tautologie, wenn sich für alle Kombinationen der Argumente eine wahre Aussage ergibt.

2.7 Beispiele

$(\mathcal{A} \Rightarrow \mathcal{B}) \Leftrightarrow (\neg \mathcal{B} \Rightarrow \neg \mathcal{A})$
denn es gilt:

$\mathcal{A}$	$\mathcal{B}$	$\mathcal{A} \Rightarrow \mathcal{B}$	$\neg \mathcal{A}$	$\neg \mathcal{B}$	$\neg \mathcal{A} \Rightarrow \neg \mathcal{B}$	$(\mathcal{A} \Rightarrow \mathcal{B}) \Leftrightarrow (\neg \mathcal{B} \Rightarrow \neg \mathcal{A})$
1	1	1	0	0	1	1
1	0	0	0	1	1	1
0	1	1	1	0	0	1
0	0	1	1	1	1	1

2.8 Satz: (wichtige Tautologien)

- (a) $\mathcal{A} \vee \neg \mathcal{A}$ Satz vom ausgeschlossenen Dritten
- (b) $\neg(\mathcal{A} \wedge \neg \mathcal{A})$ Satz vom Widerspruch
- (c) $\neg(\neg \mathcal{A}) \Leftrightarrow \mathcal{A}$ doppelte Verneinung
- (d) $\left. \begin{array}{l} (\mathcal{A} \wedge \mathcal{B}) \Leftrightarrow (\mathcal{B} \wedge \mathcal{A}) \\ (\mathcal{A} \vee \mathcal{B}) \Leftrightarrow (\mathcal{B} \vee \mathcal{A}) \end{array} \right\}$ Kommutativgesetz
- (e) $\left. \begin{array}{l} \mathcal{A} \wedge (\mathcal{B} \vee \mathcal{C}) \Leftrightarrow (\mathcal{A} \wedge \mathcal{B}) \vee (\mathcal{A} \wedge \mathcal{C}) \\ \mathcal{A} \vee (\mathcal{B} \wedge \mathcal{C}) \Leftrightarrow (\mathcal{A} \vee \mathcal{B}) \wedge (\mathcal{A} \vee \mathcal{C}) \end{array} \right\}$ Distributivgesetz
- (f) $\left. \begin{array}{l} \neg(\mathcal{A} \wedge \mathcal{B}) \Leftrightarrow (\neg \mathcal{A}) \vee (\neg \mathcal{B}) \\ \neg(\mathcal{A} \vee \mathcal{B}) \Leftrightarrow (\neg \mathcal{A}) \wedge (\neg \mathcal{B}) \end{array} \right\}$ de Morgan'sche Gesetze
- (g) $(\mathcal{A} \Rightarrow \mathcal{B}) \Leftrightarrow (\neg \mathcal{B} \Rightarrow \neg \mathcal{A})$ Kontraposition (siehe 2.7)
- (h) $(\mathcal{A} \Rightarrow \mathcal{B}) \Leftrightarrow \neg \mathcal{A} \vee \mathcal{B}$
- (i) $\left. \begin{array}{l} ((\mathcal{A} \wedge \mathcal{B}) \wedge \mathcal{C}) \Leftrightarrow (\mathcal{A} \wedge (\mathcal{B} \wedge \mathcal{C})) \\ ((\mathcal{A} \vee \mathcal{B}) \vee \mathcal{C}) \Leftrightarrow (\mathcal{A} \vee (\mathcal{B} \vee \mathcal{C})) \end{array} \right\}$ Assoziativgesetze
- (j) $\left. \begin{array}{l} (\mathcal{A} \vee \mathcal{A}) \Leftrightarrow \mathcal{A} \\ (\mathcal{A} \wedge \mathcal{A}) \Leftrightarrow \mathcal{A} \end{array} \right\}$ Idempotenz

2.9 Beispiele

Mit Satz 2.8 zeigen wir

$$((\mathcal{A} \Rightarrow \mathcal{B}) \wedge (\mathcal{B} \Rightarrow \mathcal{C})) \Rightarrow (\mathcal{A} \Rightarrow \mathcal{C})$$

Beweis:

$$(((\mathcal{A} \Rightarrow \mathcal{B}) \wedge (\mathcal{B} \Rightarrow \mathcal{C})) \Rightarrow (\mathcal{A} \Rightarrow \mathcal{C}))$$

$$\begin{aligned}
 &\stackrel{(h)}{\Leftrightarrow} (\neg((\mathcal{A} \Rightarrow \mathcal{B}) \wedge (\mathcal{B} \Rightarrow \mathcal{C})) \vee (\mathcal{A} \Rightarrow \mathcal{C})) \\
 &\stackrel{(f)}{\Leftrightarrow} ((\neg(\mathcal{A} \Rightarrow \mathcal{B}) \vee \neg(\mathcal{B} \Rightarrow \mathcal{C})) \vee (\mathcal{A} \Rightarrow \mathcal{C})) \\
 &\stackrel{(h)}{\Leftrightarrow} ((\neg(\neg\mathcal{A} \vee \mathcal{B}) \vee \neg(\neg\mathcal{B} \vee \mathcal{C})) \vee \neg(\mathcal{A} \vee \mathcal{C})) \\
 &\stackrel{(f),(c)}{\Leftrightarrow} (\mathcal{A} \wedge \neg\mathcal{B}) \vee (\mathcal{B} \wedge \neg\mathcal{C}) \vee (\neg\mathcal{A} \vee \mathcal{C}) \\
 &\stackrel{(d)}{\Leftrightarrow} ((\mathcal{A} \wedge \neg\mathcal{B}) \vee \neg\mathcal{A} \vee (\mathcal{B} \wedge \neg\mathcal{C}) \vee \mathcal{C}) \\
 &\stackrel{(e)}{\Leftrightarrow} (\underbrace{(\mathcal{A} \vee \neg\mathcal{A})}_1 \wedge (\neg\mathcal{B} \vee \neg\mathcal{A}) \vee (\mathcal{B} \vee \mathcal{C}) \wedge \underbrace{(\neg\mathcal{C} \vee \mathcal{C})}_1) \\
 &\stackrel{(a)}{\Leftrightarrow} 1 \wedge (\neg\mathcal{B} \vee \neg\mathcal{A}) \vee (\mathcal{B} \vee \mathcal{C}) \wedge 1 \\
 &\Leftrightarrow (\neg\mathcal{B} \vee \neg\mathcal{A}) \vee (\mathcal{B} \vee \mathcal{C}) \\
 &\stackrel{(i),(d)}{\Leftrightarrow} \underbrace{(\neg\mathcal{B} \vee \mathcal{B})}_1 \vee \neg\mathcal{A} \vee \mathcal{C} \\
 &\Leftrightarrow 1 \vee \neg\mathcal{A} \vee \mathcal{C} \\
 &\Leftrightarrow 1 \quad \checkmark
 \end{aligned}$$

□

2.10 Quantoren

dienen der kompakten Schreibweise logischer Ausdrücke:

- $\forall$: für alle
- $\exists$: es existiert ein
ebenfalls gebräuchlich:
- $\exists!$: es existiert genau ein
- $\nexists$: es existiert kein

Beispiel: $\forall x \mathcal{A}(x)$: für alle x gilt $\mathcal{A}(x)$.

2.11 Negation von Aussagen

Negation vertauscht Quantoren $\forall$ und $\exists$ und sie negiert die Aussage:

$$\begin{aligned}\neg(\exists x \mathcal{A}(x)) &\Leftrightarrow \forall x (\neg \mathcal{A}(x)) \\ \neg(\forall x \mathcal{A}(x)) &\Leftrightarrow \exists x (\neg \mathcal{A}(x))\end{aligned}$$

Beispiel: „Alle Radfahrer haben krumme Beine“.

Negation: „Es gibt einen Radfahrer der keine krummen Beine hat.“

Bei komplizierteren Ausdrücken geht man sukzessive vor:

$$\begin{aligned}&\neg(\exists y \forall x \mathcal{A}(x, y)) \\ \Leftrightarrow &\forall y \neg(\forall x \mathcal{A}(x, y)) \\ \Leftrightarrow &\forall y \exists x \neg \mathcal{A}(x, y)\end{aligned}$$

Kapitel 3

Beweisprinzipien

3.1 Bedeutung in der Informatik

Informatiker beweisen ständig, sie nennen es oftmals nur nicht so.

- „Tut ein Protokoll, was es soll?“
- „Arbeitet ein Algorithmus in allen Spezialfällen richtig?“
- „Bricht er in endlicher Zeit ab?“

Einige der Tautologien aus Satz 2.8 auf Seite 31 ermöglichen entsprechende Beweisprinzipien.

3.2 Direkter Beweis

- Man möchte $\mathcal{A} \Rightarrow \mathcal{B}$ zeigen.
- Verwende Tautologie (vgl. 2.9)
$$((A \Rightarrow C_1) \wedge (C_1 \Rightarrow C_2) \wedge (C_2 \Rightarrow C_3) \wedge \dots \wedge (C_k \Rightarrow C_{k+1}) \wedge (C_{k+1} \Rightarrow B)) \Leftrightarrow (A \Rightarrow B)$$
- Man zeigt also $\mathcal{A} \Rightarrow \mathcal{B}$ in einer Kette von einfacheren Implikationen.

3.3 Beispiel:

Satz: Ist eine natürliche Zahl durch 6 teilbar, so ist sie auch durch 3 teilbar.

Beweis: Sei $n \in \mathbb{N}$ durch 6 teilbar.

$$\begin{aligned} \exists k \in \mathbb{N} : n &= 6k \\ \Rightarrow 3 \cdot \underbrace{2k}_p &= 3p \text{ mit } p = 2k \in \mathbb{N} \end{aligned}$$

$\Rightarrow n$ ist durch 3 teilbar.

□

3.4 Beweis durch Kontraposition

- Beruht auf Tautologie $(\mathcal{A} \Rightarrow \mathcal{B}) \Leftrightarrow (\neg \mathcal{B} \Rightarrow \neg \mathcal{A})$.
- Statt $\mathcal{A} \Rightarrow \mathcal{B}$ zeigt man also $\neg \mathcal{B} \Rightarrow \neg \mathcal{A}$.

3.5 Widerspruchsbeweise

- Um Aussage $\mathcal{B}$ zu beweisen, zeigt man z.B., dass die Negation $\neg \mathcal{B}$ zu einem Widerspruch führt.
 $\neg \mathcal{B} \Rightarrow \mathcal{B}$
- Also muss $\neg(\neg \mathcal{B})$ und damit $\mathcal{B}$ wahr sein.
- beruht auf Tautologie $(\neg \mathcal{B} \Rightarrow \neg \mathcal{B}) \Rightarrow \mathcal{B}$.

3.6 Beispiel:

Satz: Es gibt keine größte Primzahl. (Aussage $\mathcal{B}$)

Beweis: Annahme: Es gibt eine größte Primzahl. (Aussage $\neg \mathcal{B}$)

Seien $p_1, p_2, \dots, p_n$ alle Primzahlen (*)

Sei $q = p_1 \cdot p_2 \cdot \dots \cdot p_n + 1$

Beh.: q ist eine Primzahl (Aussage $\mathcal{A}$).

Anmerkung: q ist keine Primzahl ($\neg \mathcal{A}$)

$$\Rightarrow q \text{ hat Primteiler} \quad p_i \in \{p_1, p_2, \dots, p_n\}$$

$$\Rightarrow \exists \alpha, \beta \in \mathbb{Z} : p_i \cdot \alpha = q = \underbrace{p_1 \cdot p_2 \cdot \dots \cdot p_n}_{p_i \cdot \beta} + 1 = \boxed{p_i \cdot \beta + 1}$$

$\Rightarrow p_i(\alpha - \beta) = 1$, wobei $\alpha - \beta$ ganze Zahl ist.

$\Rightarrow p_i = 1$ zu p_i Primteiler. (also gilt (*))

Da q eine Primzahl ist und $q > p_i \quad \forall i \in \{1, \dots, n\}$ ist dies ein Widerspruch.

□

3.7 Äquivalenzbeweise

- beruhen auf Tautologie $(\mathcal{A} \Leftrightarrow \mathcal{B}) \Leftrightarrow ((\mathcal{A} \Rightarrow \mathcal{B}) \wedge (\mathcal{B} \Rightarrow \mathcal{A}))$.
- Um $\mathcal{A} \Leftrightarrow \mathcal{B}$ zu zeigen, beweist man also $\mathcal{A} \Rightarrow \mathcal{B}$ und $\mathcal{B} \Rightarrow \mathcal{A}$.
- Will man die Äquivalenz vieler Aussagen zeigen, bietet sich ein Ringschluss an:
 $\mathcal{A} \Leftrightarrow \mathcal{B} \Leftrightarrow \mathcal{C} \Leftrightarrow \mathcal{D}$ zeigt man durch $\mathcal{A} \Rightarrow \mathcal{B}, \mathcal{B} \Rightarrow \mathcal{C}, \mathcal{C} \Rightarrow \mathcal{D}, \mathcal{D} \Rightarrow \mathcal{A}$.

3.8 Beweis durch vollständig Induktion

Grundidee:

- Für jede natürliche Zahl $n \in \mathbb{N}$ sei eine Aussage $A(n)$ gegeben.
- Es gelte:
 1. Induktionsanfang: $\mathcal{A}(1)$ ist wahr.
 2. Induktionsschluss: $(\mathcal{A}(n) \Rightarrow \mathcal{A}(n+1))$ ist wahr.
- Dann gilt die Aussage für alle $n \in \mathbb{N}$.

3.9 Beispiel:

Definition:

$$\sum_{k=m}^n a_k = a_m + a_{m+1} + \dots + a_n \quad \text{Summenzeichen}$$
$$\prod_{k=m}^n a_k = a_m \cdot a_{m+1} \cdot \dots \cdot a_n \quad \text{Produktzeichen}$$

Satz:

$$\sum_{k=1}^n k = 1 + 2 + 3 + \dots + n = \frac{n(n+1)}{2} \quad \text{für alle } n \in \mathbb{N}$$

Beweis mit vollständiger Induktion

1. Induktionsanfang:

$$\text{Für } n = 1 \text{ ist } \sum_{k=1}^1 k = 1 = \frac{1(1+1)}{2} = 1 \quad \text{OK.}$$

2. Induktionsschluss:

Annahme: $\mathcal{A}(n)$ für ein bestimmtes $n \in \mathbb{N}$ wahr (Ind.-Voraussetzung)

$$\sum_{k=1}^n k = \frac{n(n+1)}{2} \quad (*)$$

Dann gilt:

$$\begin{aligned} \sum_{k=1}^{n+1} k &= n+1 + \sum_{k=1}^n k \\ &\stackrel{(*)}{=} \frac{n(n+1)}{2} + (n+1) \\ &= \frac{2(n+1) + n(n+1)}{2} \\ &= \frac{(n+1)(n+2)}{2} \end{aligned}$$

d.h. $\mathcal{A}(n+1)$ ist wahr.

□

3.10 Anmerkung:

- (a) Induktionsbeweise sind sehr häufig bei Summen- und Produktformeln. Diese sind u.a. wichtig für Komplexitätsabschätzungen von Algorithmen.
- (b) Der Induktionsanfang muss nicht bei 1 beginnen. Beginnt er mit k so gilt die Aussage $\mathcal{A}(n)$ für alle $n \in \mathbb{N}$ mit $n \geq k$.

Kapitel 4

Relationen

4.1 Motivation

Wir möchten:

- die Elemente zweier Mengen in Beziehung setzen.
- eine Menge so unterteilen, dass „ähnliche“ Elemente in der selben Klasse liegen.
- die Elemente innerhalb einer Klasse ordnen.

In der Informatik sind solche Fragen, z.B. bei relationalen Datenbanken wichtig.

4.2 Definition: *(Relation)*

Seien A, B nichtleere Mengen.

Eine Relation auf $A \times B$ ist eine Teilmenge $R \subset A \times B = \{(x, y) \mid x \in A, y \in B\}$.

Für $(x, y) \in R$ sagt man: „ x steht in Relation R zu y “. Man schreibt auch: xRy .

4.3 Beispiel:

- (a) $R_1 = \{(x, y) \in \mathbb{R} \mid x = y^2\}$
 $R_2 = \{(x, y) \in \mathbb{R} \mid x < y\}$
sind Relationen auf $\mathbb{R}^2$

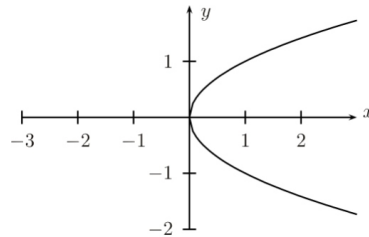


Abbildung 4.1: $\{(x, y) \in \mathbb{R} \mid x = y^2\}$

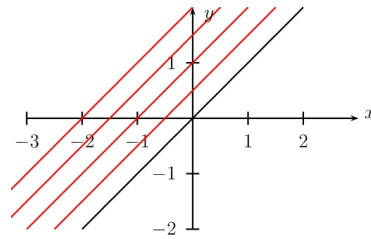


Abbildung 4.2: $\{(x, y) \in \mathbb{R} \mid x < y\}$

- (b) Sei S Menge der Studierenden der Universität des Saarlandes und F sei die Menge der Studiengänge. Dann ist
 $B := \{(s, f) \mid s \text{ belegt den Studiengang } f\}$
eine Relation auf $S \times F$. In relationalen Datenbanken werden Datensätze durch derartige Relationen charakterisiert.

4.4 Reflexivität, Symmetrie, Transitivität bei Relation

Eine Relation $R \subset A \times A$ heißt:

- (a) reflexiv, falls $xRx \quad \forall x \in A$
(b) symmetrisch, falls $xRy \Rightarrow yRx \quad \forall x, y \in A$
(c) transitiv, falls aus xRy und yRz stets xRz folgt.

Ist eine Relation R reflexiv, symmetrisch und transitiv so heißt sie Äquivalenzrelation.

Man schreibt auch $x \sim y$ statt xRy .

4.5 Beispiele

- (a) Sei A die Menge der Einwohner von Deutschland. Wir definieren die Relation
 $xRy: \Leftrightarrow x$ und y haben ihren ersten Wohnsitz in der selben Stadt.
Dann ist R eine Äquivalenzrelation.

(b) $A := \mathbb{Z}$. Betrachte die Relation

$$R_5 = \{(x, y) \in \mathbb{Z}^2 \mid x - y \text{ ist ohne Rest durch } 5 \text{ teilbar}\}.$$

Dann gilt:

(i) $xR_5x \quad \forall x \in \mathbb{Z}$, da $x-x = 0$ durch 5 teilbar ist. (Reflexivität)

(ii) Sei xR_5y :

$$\Rightarrow \exists k \in \mathbb{Z} : x - y = 5k$$

$$\Rightarrow y - x = 5 \cdot (-k)$$

$$\Rightarrow yR_5x \quad (\text{Symmetrie})$$

(iii) Sei xR_5y und yR_5z .

$$\Rightarrow \exists k, l \in \mathbb{Z} :$$

$$x - y = 5k$$

$$y - z = 5l$$

$$\Rightarrow x - z = (x - y) + (y - z)$$

$$= 5k - 5l$$

$$= 5(k - l)$$

$$\Rightarrow xR_5z \quad (\text{Transitivität})$$

Damit ist R_5 eine Äquivalenzrelation.

(c) $A := \mathbb{R}$

$$R := \{(x, y) \in \mathbb{R} \mid x < y\}.$$

R ist keine Äquivalenzrelation, da nicht reflexiv und nicht symmetrisch, aber transitiv.

Äquivalenzrelation erlauben eine Einteilung in Klassen:

4.6 Definition: (Äquivalenzklassen)

Sei $\sim$ eine Äquivalenzrelation auf A und $a \in A$.

Dann heißt $[a] := \{x \in A \mid x \sim a\}$ die Äquivalenzklasse von a . Die Elemente von $[a]$ heißen die zu a äquivalenten Elemente.

4.7 Beispiel:

Sei R_5 definiert wie in 4.5(b). Dann ist $[a] := \{x \in \mathbb{Z} \mid x - a \text{ ist durch } 5 \text{ teilbar}\}$. Insbesondere gilt:

$$\bullet [0] := \{0, 5, 10, 15, \dots, -5, -10, -15, \dots\}$$

$$\bullet [1] := \{1, 6, 11, 16, \dots, -4, -9, -14, \dots\}$$

- $[2] := \{2, 7, 12, 17, \dots, -3, -8, -13, \dots\}$
- $[3] := \{3, 8, 13, 18, \dots, -2, -7, -12, \dots\}$
- $[4] := \{4, 9, 14, 19, \dots, -1, -6, -11, \dots\}$

Offenbar gilt: $[5] = [0]$, $[6] = [1]$, $[7] = [2]$, ...

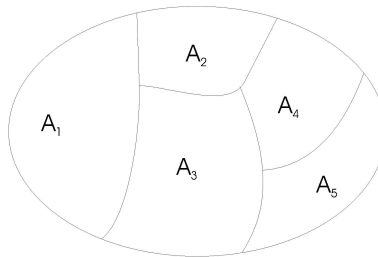
Man nennt $\mathbb{Z}_5 := \{[0], [1], \dots, [4]\}$ die Restklassen von $\mathbb{Z}$ modulo 5.

Es gilt: $\mathbb{Z} = [0] \cup [1] \cup [2] \cup [3] \cup [4]$ und ferner sind die Äquivalenzklassen $[0], [1], \dots, [4]$ disjunkt. Ist dies Zufall?

4.8 Definition: (Partition)

Eine Partition einer Menge A ist eine Menge $P = \{A_1, A_2, \dots\}$ von nichtleeren Teilmengen von A mit

- (a) $A_i \cap A_j \neq \emptyset$ für $i \neq j$
- (b) $\bigcup_i A_i = A$.



Man schreibt auch: $A = \dot{\bigcup}_i A_i$.

4.9 Satz: (Partitionseigenschaften von Äquivalenzklassen)

Sei $\sim$ eine Äquivalenzrelation auf M . Dann bildet die Menge der Äquivalenzklassen eine Partition von M .

Beweis:

- (a) Seien A, B Äquivalenzklassen. Durch Kontraposition zeigen wir: Falls $A \neq B$, sind A, B disjunkt.
Seien also A, B nicht disjunkt: $A \cap B \neq \emptyset$.
Seien $[a] = A, [b] = B$ und $y \in A \cap B$. Um $A = B$ zu zeigen, beweisen wir $A \subset B$ und $B \subset A$:

„ $A \subset B$ “:

Sei $x \in A \Rightarrow x \sim a$
 Da $y \in A \Rightarrow y \sim a \xRightarrow{\text{Symmetrie}} a \sim y$ } $\xRightarrow{\text{Transitivitt}} x \sim y$ } $\Rightarrow x \sim b,$
 Wegen $y \in B : y \sim b$ d.h. $x \in B$
 „ $B \subset A$ “: analog

- (b) Für jedes $x \in A$ ist $x \in [x]$ (wegen Reflexivität)
 $\Rightarrow A = \bigcup_{x \in A} [x].$

□

Eine Äquivalenzrelation verhält sich im Prinzip wie $=$. Wie lassen sich Relationen charakterisieren, die sich im Prinzip wie $\leq$ verhalten, d.h. als Ordnungsrelationen wirken?

4.10 Definition: (Teilrelation)

Sei $A \neq \emptyset$. Eine Relation $R \subset A \times A$ heißt Teilordnung auf A , wenn gilt:

- (a) R ist reflexiv: $xRx \quad \forall x \in A$
 (b) R ist transitiv: $xRy, yRx \Rightarrow xRz \quad \forall x, y, z \in A$
 (c) R ist antisymmetrisch $xRy, yRx \Rightarrow x = y \quad \forall x, y \in A$

Ferner heißt R total, wenn xRy oder yRx für alle $x, y \in R$ gilt (d.h. wenn x und y vergleichbar sind). Eine totale Teilordnung nennen wir auch (Total-) Ordnung.

4.11 Beispiele

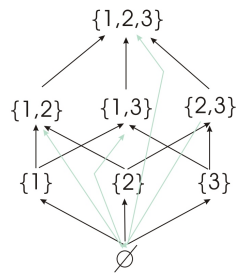
- (a) Die Relation „ $\leq$ “ definiert eine Teilordnung auf $\mathbb{R}$:

- (i) $x \leq x \quad \forall x \in \mathbb{R}$ (Reflexivität)
 (ii) $x \leq y, y \leq z \Rightarrow x \leq z \quad \forall x, y, z \in \mathbb{R}$ (Symmetrie)
 (iii) $x \leq y, y \leq x \Rightarrow x = y \quad \forall x, y \in \mathbb{R}$ (Transitivität)

Da $x \leq y$ oder $y \leq x$ für alle $x, y \in \mathbb{R}$, ist $\leq$ eine (totale) Ordnung.

- (b) Sei M eine Menge. Betrachte die Relation „ $\subset$ “ auf der Potenzmenge $P(M)$. Dann ist „ $\subset$ “ eine Teilordnung

- (i) $A \subset A \quad \forall A \in P(M)$ (Reflexivität)
 (ii) $A \subset B, B \subset C \Rightarrow A \subset C \quad \forall A, B, C \in P(M)$ (Transitivität)
 (iii) $A \subset B, B \subset A \Rightarrow A = B \quad \forall A, B \in P(M)$ (Antisymmetrie)



Beispiel: : $M = \{1, 2, 3\}$

Dabei symbolisiert $\rightarrow$ die Relation $\subset$. Allerdings ist $\subset$ keine Totalordnung, da es nicht vergleichbare Elemente in $P(M)$ gibt.

z.B. sind $\{2\}$ und $\{1, 3\}$ nicht vergleichbar sind.

- (c) Sei M die Menger der englischen Wörter mit 4 Buchstaben. Verwendet man die alphabetische Ordnung und vereinbart, dass Großbuchstaben vor Kleinbuchstaben kommen, liegt eine Totalordnung vor (lexikographische Ordnung).

Kapitel 5

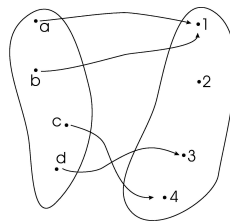
Abbildungen

5.1 Definition: (Abbildung)

Eine Abbildung zwischen zwei Mengen M und N ist eine Vorschrift $f : M \rightarrow N$, die jedem Element $x \in M$ ein Element $f(x) \in N$ zuordnet.
Schreibweise: $x \mapsto f(x)$

5.2 Beispiel:

$M = \{a, b, c, d\}, N = \{1, 2, 3, 4\}$
 $a \mapsto f(a) = 1$



5.3 Wichtige Begriffe

Für eine Teilmenge $A \subset M$ heißt $f(A) := \{f(x) \mid x \in A\}$ das Bild von A .

Für $B \subset N$ heißt $f^{-1}(B) := \{x \in M \mid f(x) \in B\}$ das Urbild von B .

Im Beispiel 5.2:

- $f(\{a, b\}) = \{1\}$,

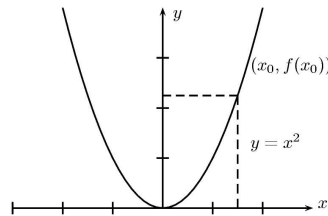
- $f(\{a, c\}) = \{1, 4\}$,
- $f^{-1}(\{3\}) = \{d\}$,
- $f^{-1}(\{2\}) = \emptyset$.

Hat B nur ein Element, $B = \{y\}$, dann setzt man $f^{-1}(y) := f^{-1}(\{y\})$.

Ist $f : M \rightarrow N$ eine Abbildung und $A \subset M$, dann heißt
 $f|_A : A \rightarrow N, A \ni x \mapsto f(x)$ die Einschränkung von f auf A .

Die Relation $\Gamma_f := \{(x, y) \in M \times N \mid y = f(x)\}$ heißt Graph von f .

Beispiel: $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$

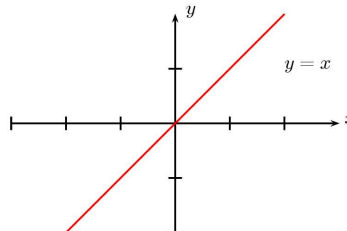


5.4 Definitionsbereich

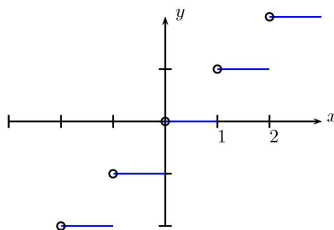
Zu den wichtigsten Abbildungen zählen reelle Funktionen $f : D \rightarrow \mathbb{R}$ mit einer nichtleeren Teilmenge $D \subset \mathbb{R}$, dem Definitionsbereich.

5.5 Beispiel:

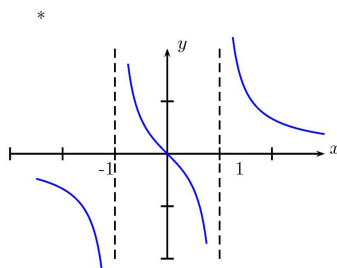
- (a) $id_{\mathbb{R}} : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x$
identische Abbildung



- (b) entier-Funktion:
entier: $\mathbb{R} \rightarrow \mathbb{Z} \subset \mathbb{R}$
entier(x) := größte ganze Zahl $\leq x$



(c) $f : D \rightarrow \mathbb{R}, x \mapsto \frac{x}{x^2-1}$ mit $D = \mathbb{R} \setminus \{-1, 1\}$.



5.6 Definition: (surjektiv, injektiv, bijektiv)

Eine Abbildung $f : M \rightarrow N$ heißt

- surjektiv, wenn es für alle $y \in N$ ein $x \in M$ gibt mit $f(x) = y$.
(d.h. $f(M) = N$, „Abbildung auf N “)
- injektiv (eindeutig), wenn keine zwei verschiedenen Elemente von M auf das selbe Element von N abgebildet werden:
 $f(x_1) = f(x_2) \Rightarrow x_1 = x_2$
- bijektiv, wenn sie injektiv und surjektiv ist.

5.7 Beispiele

- (a) $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$ ist
- nicht surjektiv: z.B. hat -1 kein Urbild
 - nicht injektiv: z.B. ist $f(2) = 4 = f(-2)$
- (b) $f : \mathbb{R} \rightarrow \mathbb{R}_0^+ := \{x \in \mathbb{R} \mid x \geq 0\}, x \mapsto x^2$ ist surjektiv, aber nicht injektiv.
- (c) $f : \mathbb{R}_0^+ \rightarrow \mathbb{R}, x \mapsto x^2$ ist injektiv, aber nicht surjektiv.
- (d) $f : \mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+, x \mapsto x^2$ ist surjektiv und injektiv, also bijektiv.

Bijektive Abbildungen kann man umkehren:

5.8 Definition: (Umkehrabbildung)

Ist eine Abbildung $f : M \rightarrow N$ bijektiv, so heißt die Abbildung

$$f^{-1} : N \rightarrow M, y \mapsto x$$

mit $y = f(x)$ die Umkehrabbildung von f .

5.9 Beispiele

$f : \mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+, x \mapsto x^2$ hat die Umkehrabbildung $f^{-1} : \mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+, x \mapsto \sqrt{x}$

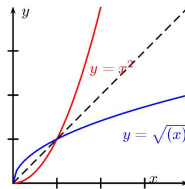


Abbildung lassen sich miteinander verknüpfen:

5.10 Definition: (Komposition)

Sind $f : A \rightarrow B$ und $g : B \rightarrow C$ Abbildungen. Dann heißt die Abbildung

$$g \circ f : A \rightarrow C, x \mapsto g(f(x))$$

die Komposition (Verknüpfung, Hintereinanderschaltung) von f und g .
Sprechweise: „ g Kringel f “.

5.11 Satz: (Assoziativität der Verknüpfung)

Seien $f : A \rightarrow B$, $g : B \rightarrow C$, $h : C \rightarrow D$ Abbildungen, so gilt:

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

d.h. die Komposition ist assoziativ.

Beweis:

Ist $x \in A$, so gilt

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

□

5.12 Vorsicht

Die Komposition von Abbildungen ist i.A. nicht kommutativ!

Beispiel: $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x + 1$

$g : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$

$\Rightarrow$

$$\begin{aligned}(f \circ g)(x) &= f(x^2) = x^2 + 1 \\ (g \circ f)(x) &= g(x + 1) = (x + 1)^2 = x^2 + 2x + 1\end{aligned}$$

d.h. $f \circ g \neq g \circ f$.

Die Komposition kann auch zum Nachweis von Injektivität, Surjektivität und Bijektivität dienen:

5.13 **Satz:** (Kriterium für Injektivität, Surjektivität, Bijektivität)

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

- (a) f injektiv $\Leftrightarrow \exists g : Y \rightarrow X$ mit $g \circ f = id_X$
- (b) f surjektiv $\Leftrightarrow \exists g : Y \rightarrow X$ mit $f \circ g = id_Y$
- (c) f bijektiv $\Leftrightarrow \exists g : Y \rightarrow X$ mit $g \circ f = id_X$ und $f \circ g = id_Y$

In diesem Fall ist g die Umkehrabbildung f^{-1} .

Beweis:

- (a) „ $\Rightarrow$ “: Sei f injektiv. Dann existiert zu jeden $y \in f(X)$ genau ein $x \in X$ mit $f(x) = y$. Setze $g(y) := x$. Für alle $y \in Y \setminus f(X)$ setzen wir $g(y) := x_0$ mit $x_0 \in X$ beliebig. Dann hat $g : Y \rightarrow X$ die Eigenschaft $g \circ f = id_X$.
„ $\Leftarrow$ “: Sei $g : Y \rightarrow X$ mit $g \circ f = id_X$ gegeben. Sei ferner $f(x_1) = f(x_2)$.
 $\Rightarrow x_1 = g(f(x_1)) = g(f(x_2)) = x_2$
 $\Rightarrow f$ ist injektiv.

- (b) „ $\Rightarrow$ “: Sei f surjektiv. Zu jedem $y \in Y$ wählen wir ein festes $x \in X$ mit $f(x) = y$ und setzen $g(y) := x$. Dann hat $g : Y \rightarrow X$ die Eigenschaft $f \circ g = id_Y$.
- „ $\Leftarrow$ “: Sei $g : Y \rightarrow X$ mit $f \circ g = id_Y$ gegeben. Sei $y \in Y$.
 $\Rightarrow y = f(g(y))$, d.h. $y \in f(X) \Rightarrow f$ ist surjektiv.
- (c) „ $\Rightarrow$ “: Sei f bijektiv. Dann existiert nach (a) und (b) eine Abbildung $g : Y \rightarrow X$ mit $g \circ f = id_X$ und $f \circ g = id_Y$.
- „ $\Leftarrow$ “: Sei $g : Y \rightarrow X$ mit $g \circ f = id_X$ und $f \circ g = id_Y$ gegeben. Dann ist f nach (a) und (b) injektiv und surjektiv, d.h. bijektiv.

Ferner erfüllt g die Def. 5.8 auf Seite 48 der Umkehrabbildung von $f : g = f^{-1}$.

□

5.14 Mächtigkeit von Mengen

Eine Menge M heißt endlich, falls ein $n \in \mathbb{N}_0$ existiert und eine bijektive Abbildung $f : \{1, 2, 3, \dots, n\} \rightarrow M$. n ist eindeutig bestimmt und heißt die Mächtigkeit (Kardinalzahl, Kardinalität) von M .

Schreibweise: $|M| = n$.

Man kann die Elemente von M aufzählen: $M = \{m_1, m_2, \dots, m_n\}$, wobei $f(i) = m_i$ für $i \in \{1, \dots, n\}$. Ist M nicht endlich, so heißt M unendlich. Zwei Mengen sind gleichmächtig, falls eine bijektive Abbildung zwischen ihnen existiert.

Wir nennen eine Menge M abzählbar, wenn eine bijektive Abbildung $f : \mathbb{N} \rightarrow M$ existiert. Als Symbol für die Kardinalzahl $|\mathbb{N}|$ benutzen wir $\aleph_0$ („aleph 0“).

5.15 Beispiele

- (a) $|\{5, 7, 2, 9\}| = 4$
- (b) $|\mathbb{Z}| = \aleph_0$, denn mit $\mathbb{Z} = \{0, 1, -1, 2, -2, \dots\}$ ergibt sich eine Aufzählung mit der bijektiven Abbildung

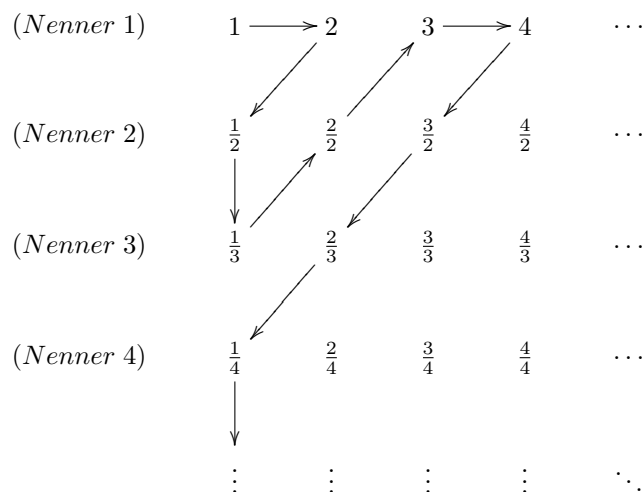
$$f : \mathbb{N} \rightarrow \mathbb{Z}, \quad f(n) = \begin{cases} \frac{n}{2} & \text{für gerades } n \\ -\frac{n-1}{2} & \text{für ungerades } n \end{cases}$$

5.16 Abzählbarkeit von $\mathbb{Q}$

$\mathbb{Q}$ ist abzählbar: $|\mathbb{Q}| = \aleph_0$.

Beweis:

(Cantor) Wir beschränken uns zunächst auf $\mathbb{Q}^+ := \{q \in \mathbb{Q} \mid q > 0\}$. Die Elemente von $\mathbb{Q}^+$ lassen sich nach folgendem Schema anordnen:



Durch Weglassen doppelter Elemente in diesem Polygonzug entsteht die Anordnung

$$1, 2, \frac{1}{2}, \frac{1}{3}, 3, 4, \frac{3}{2}, \frac{2}{3}, \frac{1}{4}, \dots$$

die sich bijektiv auf $\mathbb{N}$ abbilden lässt.

Analog wie in 5.15 auf der vorherigen Seite(b) dehnt man den Beweis auf ganz $\mathbb{Q}$ aus.

□

Bemerkung: Nicht alle unendliche Mengen sind gleichmächtig. Z.B. ist $\mathbb{R}$ überabzählbar (d.h. nicht abzählbar).

5.17 Satz: (Äquivalenz von Surjektivität und Injektivität)

Sei $f : X \rightarrow Y$ eine Abbildung zwischen endlichen, gleichmächtigen Mengen X und Y . Dann sind äquivalent:

- (a) f ist injektiv.
- (b) f ist surjektiv.
- (c) f ist bijektiv.

Beweis:

Wegen der Def. der Bijektivität genügt es, $(a) \Leftrightarrow (b)$ zu zeigen.

„ $(a) \Rightarrow (b)$ “: Sei $f : X \rightarrow Y$ injektiv.

$$\Rightarrow |f^{-1}(y)| \leq 1 \quad \forall y \in Y \quad (\text{einige } y \in Y \text{ könnten kein Urbild haben})$$

$$\Rightarrow |X| = \sum_{y \in Y} |f^{-1}(y)| \leq \sum_{y \in Y} 1 = |Y|$$

Wegen $|X| = |Y|$ muss $|f^{-1}(y)| = 1 \quad \forall y \in Y$ gelten.

$\Rightarrow f$ ist surjektiv, da jedes Element in Y zu $f(X)$ gehört.

„ $(b) \Rightarrow (a)$ “: Sei $f : X \rightarrow Y$ surjektiv.

$$\Rightarrow |Y| = \sum_{y \in Y} 1 \leq \sum_{y \in Y} |f^{-1}(y)| = |X|.$$

Wegen $|X| = |Y|$ muss $|f^{-1}(y)| = 1 \quad \forall y \in Y$ gelten.

$\Rightarrow f$ ist injektiv, da keine zwei Elemente aus X auf das selbe $y \in Y$ abgebildet werden.

□

5.18 Folgerungen

Aus dem Beweis von 5.17 auf der vorherigen Seite folgt für endliche Mengen X, Y :

(a) Ist $f : X \rightarrow Y$ injektiv, dann ist $|X| \leq |Y|$.

(b) Ist $f : X \rightarrow Y$ surjektiv, dann ist $|X| \geq |Y|$.

Die Kontraposition zu (a) liefert:

5.19 Schubfachprinzip

Seien X, Y endliche Mengen. Dann ist eine Abbildung $f : X \rightarrow Y$ mit $|X| > |Y|$ nicht injektiv. Anders formuliert: Seien m Objekte in n Kategorien („Schubfächern“) eingeteilt. Wenn $m > n$ ist, gibt es mindestens eine Kategorie, die mindestens 2 Objekte enthält.

5.20 Beispiele

- (a) Unter 13 Personen gibt es mindestens 2, die im selben Monat Geburtstag haben.
- (b) In jeder Gruppe mit mindestens zwei Personen gibt es zwei, die die gleiche Anzahl von Bekannten innerhalb dieser Gruppe haben.

Beweis:

Wir nehmen an, dass „bekannt“ eine symmetrische, nicht reflexive Relation ist.

Objekte: Personen der Gruppe. Annahme: m Personen

Kategorien: Personen mit der gleichen Zahl von Bekannten

$K_0, K_1, K_2, \dots, K_{m-1}$: Personen mit $0, 1, 2, \dots, m-1$ Bekannten.

Schubfachprinzip nicht direkt anwendbar, da die Objektzahl m identisch mit der Kategorienzahl ist. Es gibt jedoch eine Kategorie, die nicht auftritt: Denn: Angenommen, eine Person ist in K_{m-1} .

$\Rightarrow$ Sie kennt alle anderen.

$\Rightarrow$ Alle anderen kennen mindestens 1 Person.

$\Rightarrow K_0$ ist leer.

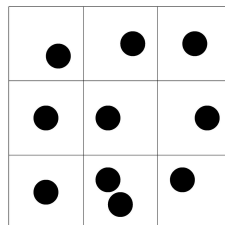
Wenn keine Person in K_{m-1} ist, ist K_{m-1} leer.

$\Rightarrow$ Also ist das Schubfachprinzip anwendbar.

□

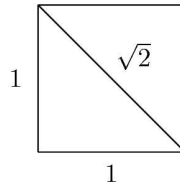
- (c) Für beliebige $n^2 + 1$ Punkte im Quadrat
 $Q := \{(x, y) \mid 0 < x < n, 0 < y < n\}$ gibt es zwei mit Abstand $\leq \sqrt{2}$

Beispiel: $n = 3$, 10 Punkte



Beweis:

Betrachte Teilquadrate $Q_{ij} := \{(x, y) \mid i-1 \leq x < i, j-1 \leq y < j\}$.
Nach dem Schufachprinzip müssen in einem der n^2 Einheitsquadrate mindestens zwei Punkte liegen. Der Maximalabstand in einem solchen Einheitsquadrat ist $\sqrt{2}$.



□

Kapitel 6

Primzahlen und Teiler

6.1 Bedeutung in der Informatik

- Wichtige Algorithmen in der Kryptographie (z.B. RSA - Algorithmus) beruhen auf grundlegenden Ergebnissen der Zahlentheorie:
Es ist leicht, zwei große Primzahlen zu multiplizieren, aber schwierig, eine große Zahl schnell in ihre Primfaktoren zu zerlegen.

6.2 Satz und Definition: *(Division mit Rest)*

Zu jeder Zahl $a \in \mathbb{Z}$ und jeder Zahl $b \in \mathbb{N}$ gibt es eindeutig bestimmte Zahlen $q, r \in \mathbb{Z}$ mit

$$a = qb + r, \quad 0 \leq r < b.$$

Wir nennen q den Quotienten und r den Rest der Division von a durch b . a heißt Dividend, b ist der Divisor.

Beweis:

Betrachte Fall $a \geq 0$. Sei q die größte ganze Zahl mit $q \cdot b \leq a$. Dann gibt es ein $r \geq 0$ mit $a = q \cdot b + r$. Ferner gilt $r < b$, denn andernfalls wäre q nicht maximal gewesen. Den Fall $a < 0$ zeigt man ähnlich.

□

Besonders interessant ist der Fall $r = 0$ und die Erweiterung $b \in \mathbb{Z} \setminus \{0\}$.

6.3 Definition: (Teilbarkeit, Teiler, Primzahl)

Seien $a, b \in \mathbb{Z}$ und $b \neq 0$. Wir sagen, b teilt a ($b \mid a$), wenn es eine ganze Zahl q gibt mit $a = q \cdot b$.

In diesem Fall heißt b Teiler von a . Falls b kein Teiler von a ist, schreibt man $b \nmid a$. Eine natürliche Zahl $p > 1$ heißt Primzahl („ist prim“), wenn sie nur die trivialen Teiler $\pm p$ und ± 1 hat. Zahlen, die nicht prim sind, heißen zusammengesetzt.

Bemerkung: In 3.6 auf Seite 36 haben wir gesehen, dass es unendlich viele Primzahlen gibt.

6.4 Beispiele

- (a) $-7 \mid 63$, denn $(-7)(-9) = 63$.
- (b) 11 ist prim.
- (c) 35 ist zusammengesetzt $35 = 5 \cdot 7$.

Folgende Teilbarkeitseigenschaften sind leicht zu zeigen:

6.5 Satz: (Teilbarkeitsregeln)

- (a) Aus $c \mid b$ und $b \mid a$ folgt $c \mid a$.
(Beispiel: $3 \mid 12$ und $12 \mid 24 \Rightarrow 3 \mid 24$)
- (b) Aus $b_1 \mid a_1$ und $b_2 \mid a_2$ folgt $b_1 b_2 \mid a_1 a_2$.
(Beispiel: $2 \mid 4$ und $7 \mid 21 \Rightarrow 14 \mid 84$)
- (c) Aus $b \mid a_1$ und $b \mid a_2$ folgt $b \mid \alpha a_1 + \beta a_2 \quad \forall \alpha, \beta \in \mathbb{Z}$.
(Beispiel: $3 \mid 6$ und $3 \mid 9 \Rightarrow 3 \mid (2 \cdot 6 + 3 \cdot 9)$)
- (d) Aus $a \mid b$ und $b \mid a$ folgt $|a| = |b|$.

6.6 Satz: (Fundamentalsatz der Zahlentheorie)

Jede natürliche Zahl $n > 1$ ist als Produkt endlich vieler (nicht notwendig verschiedener) Primzahlen darstellbar (Primzahlfaktorisation). Diese Zerlegung ist bis auf die Reihenfolge der Faktor eindeutig.

Beweis:

siehe z.B. Brill: Mathematik für Informatiker, S. 60-63

Beispiel: $84 = 2^2 \cdot 3 \cdot 7$ ist eine Primzahlfaktorisation.

6.7 Sieb des Erathostenes (Primzahlfaktorisationen)

Primzahlfaktorisationen großer Zahlen sind aufwändig. Ein einfaches (nicht sehr effizientes) Verfahren ist das „Sieb des Erathostenes“.

- Um zu prüfen, ob n prim oder zusammengesetzt ist, genügt es für jede Primzahl $p \leq \sqrt{n}$ zu testen, ob $p \mid n$.
- Findet man einen Teiler p , kann man das Verfahren mit n/p fortsetzen.

6.8 Beispiel:

Zur Primzahlfaktorisation von 84 genügt es, alle Primzahlen $\leq \sqrt{84} \approx 9,17$ zu betrachten, d.h. 2,3,5,7. Wegen $7 \mid 84$, führt man die Faktorisierung mit $84 / 7 = 12$ fort. Hier müssen nur noch die Primzahlen 2 und 3 getestet werden.

$$\begin{aligned} 12 / 3 &= 4 \\ 4 &= 2 \cdot 2 \\ \Rightarrow 84 &= 2^2 \cdot 3 \cdot 7. \end{aligned}$$

6.9 größter gemeinsamer Teiler (ggT)

Sind $a, b, d \in \mathbb{Z}$ und gilt $d \mid a$ und $d \mid b$, so heißt d gemeinsamer Teiler von a und b . Wenn für jeden anderen gemeinsamen Teiler c von a und b gilt $c \mid d$, dann heißt d größter gemeinsamer Teiler (ggT). (engl. greatest common divisor, gcd):

$$d = \text{ggT}(a, b).$$

Beispiel: $\text{ggT}(84, 66) = 6$, denn 84 und 66 haben die Primfaktorzerlegung

$$\begin{aligned} 84 &= 2^2 \cdot 3 \cdot 7 \\ 66 &= 2 \cdot 3 \cdot 11. \end{aligned}$$

ggT ist Produkt der gemeinsamen Faktoren 2 und 3.
Gibt es schnelle Algorithmen zur Bestimmung des ggT?
Hierzu benötigen wir einen Hilfsatz (Lemma).

6.10 Lemma: (Eigenschaften des ggT)

Seien $a, b, q \in \mathbb{Z}$, Dann gilt:

$$(a) \quad d = \text{ggT}(a, b) \Leftrightarrow d = \text{ggT}(b, a - q \cdot b).$$

$$(b) \quad \text{Ist } a = q \cdot b, \text{ so gilt } b = \text{ggT}(a, b).$$

Beispiel:

$$(a) \quad 6 = \text{ggT}(84, 66) \Leftrightarrow 6 = \text{ggT}(66, 84 - 1 \cdot 66) = \text{ggT}(66, 18)$$

$$(b) \quad \text{Ist } 84 = 7 \cdot 12 \Rightarrow 12 = \text{ggT}(84, 12)$$

Beweis:

(a) Wir zeigen nur „ $\Rightarrow$ “ („ $\Leftarrow$ “ geht ähnlich).

Sei $d = \text{ggT}(a, b)$.

Aus $d \mid a$ und $d \mid b$ folgt mit 6.5 auf Seite 56(e): $d \mid a - q \cdot b$.

Also ist d gemeinsamer Teiler von b und $a - q \cdot b$.

Um zu zeigen, dass d auch grösster gemeinsamer Teiler ist, betrachten wir einen weiteren gemeinsamen Teiler c und zeigen $c \mid d$.

$$\left. \begin{array}{l} c \mid b \Rightarrow c \mid qb \\ c \mid a - qb \end{array} \right\} \Rightarrow c \mid (a - qb) + qb \Rightarrow c \mid a \quad \left. \begin{array}{l} c \mid a \\ c \mid b \end{array} \right\} \Rightarrow c \mid d, \text{ da } d = \text{ggT}(a, b)$$

(b) Prüfe Def. des ggT für b nach:

Aus $a = qb$ folgt $b \mid a$. Wegen $b \mid b$ ist b gemeinsamer Teiler von a und b .

Ist c weiterer gemeinsamer Teiler von a und b gilt $c \mid a$ und $c \mid b$.

Wegen $c \mid b$ ist $b = \text{ggT}(a, b)$.

□

Dieses Lemma bildet die Grundlage des Euklidischen Algorithmus!

6.11 Satz: (Euklidischer Algorithmus)

Für $a, b \in \mathbb{N}$ mit $a > b$ setzen wir $r_0 := a, r_1 := b$ und berechnen folgende Divisionen mit Rest:

$$\begin{array}{rclcl}
 r_0 & = & q_0 r_1 & + & r_2 & (0 < r_2 < r_1) \\
 & \swarrow & & \swarrow & & \\
 r_1 & = & q_1 r_2 & + & r_3 & (0 < r_3 < r_2) \\
 & \swarrow & & \swarrow & & \\
 r_2 & = & q_2 r_3 & + & r_4 & (0 < r_4 < r_3) \\
 & & & & & \\
 & & \vdots & & & \\
 & & & & & \\
 r_{n-2} & = & q_{n-2} r_{n-1} & + & r_n & (0 < r_n < r_{n-1}) \\
 & & & & & \\
 r_{n-1} & = & q_{n-1} r_n & & &
 \end{array}$$

Dann ist $r_n = \text{ggT}(a, b)$.

6.12 Beispiel:

$\text{ggT}(133, 91)$

$$\begin{array}{rcl}
 133 & = & 1 \cdot 91 + 42 \\
 91 & = & 2 \cdot 42 + 7 \\
 42 & = & 6 \cdot 7 \\
 7 & = & \text{ggT}(133, 91)
 \end{array}$$

6.13

Warum liefert der Euklidische Algorithmus den ggT?

Beweis:

(von Satz 6.11): Die Reste $r_j > 0$ werden in jedem Schritt kleiner.
 $\Rightarrow$ Algorithmus bricht nach endlich vielen Schritten mit Rest 0 ab.

Falls $ggT(a, b) = ggT(r_0, r_1)$ existiert, gilt nach Lemma 6.10 auf Seite 58 (a)

$$\begin{aligned}
 ggT(r_0, r_1) &= ggT(r_1, \underbrace{r_0 - q_0 r_1}_{r_2}) = ggT(r_1, r_2) \\
 ggT(r_1, r_2) &= ggT(r_2, \underbrace{r_1 - q_1 r_2}_{r_3}) = ggT(r_2, r_3) \\
 &\vdots \\
 ggT(r_{n-2}, r_{n-1}) &= ggT(r_{n-1}, \underbrace{r_{n-2} - q_{n-2} r_{n-1}}_{r_n}) = ggT(r_{n-1}, r_n)
 \end{aligned}$$

und somit $ggT(a, b) = ggT(r_{n-1}, r_n)$

Da $r_{n-1} = q_{n-1} r_n$ folgt mit Lemma 6.10 auf Seite 58(b):

$$r_n = ggT(r_{n-1}, r_n)$$

und somit $r_n = ggT(a, b)$.

□

Kapitel 7

Modulare Arithmetik

7.1 Bedeutung in der Informatik

Rechnen mit Kongruenzen ist u.A. wichtig für

- Suche von Datensätzen in großen Dateien (Hashing)
- Prüfwerte (ISBN, Barcodes, ...)
- Kryptographie
- Generierung von Pseudozufallszahlen

7.2 Definition: *(kongruent modulo m , Modul)*

Zwei ganze Zahlen a, b heißen kongruent modulo m , wenn $m \in \mathbb{N}$ ein Teiler von $a - b$ ist: $m \mid a - b$. Wir schreiben $a \equiv b \pmod{m}$, und wir nennen m Modul.

Beispiel:

- (a) $7 \equiv 22 \pmod{5}$, da $5 \mid 7 - 22$
(b) $8 \not\equiv 22 \pmod{5}$, da $5 \nmid 8 - 22$.

7.3 Satz: *(Zusammenhang Kongruenz - Division mit Rest)*

Es gilt $a \equiv b \pmod{m}$ genau dann, wenn a und b bei der Division durch m den selben Rest besitzen.

Beweis:

„ $\Rightarrow$ “ Sei $a \equiv b \pmod{m}$.

Seien ferner $q_1, q_2, r_1, r_2 \in \mathbb{Z}$ mit

$$\begin{aligned}a &= q_1 m + r_1 \\b &= q_2 m + r_2 \\und\ 0 &\leq r_1, r_2 < m \\ \Rightarrow a - b &= (q_1 - q_2) \cdot m + (r_1 - r_2).\end{aligned}$$

Wegen $m \mid a - b$ ist $r_1 - r_2 = 0$, d.h. $r_1 = r_2$.

„ $\Leftarrow$ “ Sei umgekehrt

$$\begin{aligned}a &= q_1 m + r \\b &= q_2 m + r \\ \Rightarrow a - b &= (q_1 - q_2) \cdot m, \text{ d.h. } m \mid a - b.\end{aligned} \quad \square$$

7.4 Interpretation als Äquivalenzrelation

Kongruenz modulo m definiert eine Äquivalenzrelation auf $\mathbb{Z}$ (vgl. § 4 auf Seite 39), denn es gilt:

(a) Reflexivität: $a \equiv a \pmod{m}$, da $m \mid 0$.

(b) Symmetrie: Sei

$$\begin{aligned}a &\equiv b \pmod{m} \\ \Rightarrow m &\mid (a - b) \\ \Rightarrow m &\mid (b - a) \\ \Rightarrow b &\equiv a \pmod{m}\end{aligned}$$

(c) Transitivität: Seien $a \equiv b \pmod{m}$ und $b \equiv c \pmod{m}$.
 $\Rightarrow m \mid (a - b)$ und $m \mid (b - c)$.

Mit Satz 6.5 auf Seite 56(c) folgt $m \mid (a - b) + (b - c)$, d.h. $m \mid a - c$ und somit $a \equiv c \pmod{m}$.

Die Äquivalenzklassen

$$[b] := \{a \in \mathbb{Z} \mid a \equiv b \pmod{m}\}$$

mit $b \in \{0, 1, \dots, m - 1\}$ heißen Restklassen von $\mathbb{Z}$ modulo m . Wir schreiben:

$$\mathbb{Z}_m := \{[0], [1], \dots, [m - 1]\}$$

für die Menge aller Restklassen modulo m .

Können wir in der m -elementigen Menge $\mathbb{Z}_m$ ähnlich rechnen wie in $\mathbb{Z}$?

7.5 Lemma: *(Addition von Elementen zweier Restklassen)*

Seien $[a]$ und $[b]$ zwei Restklassen (modulo m), und seien $a' \in [a]$ und $b' \in [b]$ beliebig. Dann ist $a' + b' \in [a + b]$.

Beweis:

Für $a' \in [a]$ und $b' \in [b]$ existieren $p, q \in \mathbb{Z}$ mit

$$\begin{aligned} a' - a &= pm \\ b' - b &= qm \\ \Rightarrow a' + b' &= (pm + a) + (qm + b) \\ &= (p + q)m + (a + b) \\ \text{d.h. } a' + b' &\in [a + b]. \end{aligned}$$

□

Dieses Lemma motiviert:

7.6 Definition: *(modulare Addition)*

Seien $[a]$ und $[b]$ zwei Restklassen. Wir definieren die (modulare) Addition von $[a]$ und $[b]$ durch.

$$[a] + [b] := [a + b].$$

7.7 Beispiel: Additionstabeln in $\mathbb{Z}_6$

+	[0]	[1]	[2]	[3]	[4]	[5]
[0]	[0]	[1]	[2]	[3]	[4]	[5]
[1]	[1]	[2]	[3]	[4]	[5]	[0]
[2]	[2]	[3]	[4]	[5]	[0]	[1]
[3]	[3]	[4]	[5]	[0]	[1]	[2]
[4]	[4]	[5]	[0]	[1]	[2]	[3]
[5]	[5]	[0]	[1]	[2]	[3]	[4]

denn es gilt z.B. $[3] + [5] = [8] = [2]$

7.8 Satz: (*Eigenschaften der modularen Addition*)

Die Addition in $\mathbb{Z}_m$ hat folgende Eigenschaften:

- (a) Kommutativgesetz: $[a] + [b] = [b] + [a] \quad \forall [a], [b] \in \mathbb{Z}_m$
- (b) Assoziativgesetz: $([a] + [b]) + [c] = [a] + ([b] + [c]) \quad \forall [a], [b], [c] \in \mathbb{Z}_m$
- (c) $[0]$ ist neutrales Element der Addition: $[a] + [0] = [a] \quad \forall [a] \in \mathbb{Z}_m$
- (d) Inverses Element: Zu jedem $[a] \in \mathbb{Z}_m$ existiert ein $[b] \in \mathbb{Z}_m$ mit $[a] + [b] = [0]$.

Beweis:

(a) $[a] + [b] = [a + b] = [b + a] = [b] + [a]$.

(b)

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

(c) $[a] + [0] = [a + 0] = [a]$.

(d) Das inverse Element zu $[a]$ ist $[m - a]$, denn: $[a] + [m - a] = [a + m - a] = [m] = [0]$.

□

7.9 Lemma: (Multiplikation von Elementen zweier Restklassen)

Seien $[a]$ und $[b]$ zwei Restklassen (modulo m) und seien $a' \in [a]$, $b' \in [b]$ beliebig. Dann ist $a' \cdot b' \in [a \cdot b]$

Beweis:

Für $a' \in [a]$ und $b' \in [b]$ existieren $p, q \in \mathbb{Z}$ mit .

$$\begin{aligned} a' - a &= pm \\ b' - b &= qm \end{aligned}$$

$$\begin{aligned} \Rightarrow a' \cdot b' &= (pm + a) \cdot (qm + b) \\ &= pqm^2 + pmb + aqm + ab \\ &= (pqm + pb + aq)m + ab \\ \Rightarrow a'b' &\equiv ab \pmod{m} \\ &= m \mid (a'b' - ab) \end{aligned}$$

d.h. $a'b' \in [a \cdot b]$.

□

Dieses Lemma motiviert:

7.10 Definition: (modulare Multiplikation)

Seien $[a], [b]$ zwei Restklassen. Wir definieren die (modulare) Multiplikation von $[a]$ und $[b]$ durch

$$[a] \cdot [b] := [a \cdot b].$$

7.11 Beispiel: Multiplikationstabellen in $\mathbb{Z}_6$

+	[0]	[1]	[2]	[3]	[4]	[5]
[0]	[0]	[0]	[0]	[0]	[0]	[0]
[1]	[0]	[1]	[2]	[3]	[4]	[5]
[2]	[0]	[2]	[4]	[0]	[2]	[4]
[3]	[0]	[3]	[0]	[3]	[0]	[3]
[4]	[0]	[4]	[2]	[0]	[4]	[2]
[5]	[0]	[5]	[4]	[3]	[2]	[1]

denn es gilt z.B. $[2] \cdot [5] = [10] = [4]$.

Beachte:

- (a) Es kommen i.A. nicht alle Elemente in jeder Zeile / Spalte vor. Manche fehlen, andere treten mehrfach auf.
- (b) Aus $[a] \cdot [b] = [0]$ folgt i.A. nicht $[a] = [0]$ oder $[b] = [0]$.

7.12 Satz: (*Eigenschaften der modularen Multiplikation*)

Die Multiplikation in $\mathbb{Z}_m$ hat folgende Eigenschaften.

- (a) Kommutativgesetz: $[a] \cdot [b] = [b] \cdot [a] \quad \forall [a], [b] \in \mathbb{Z}_m$
- (b) Assoziativgesetz: $([a] \cdot [b]) \cdot [c] = [a] \cdot ([b] \cdot [c]) \quad \forall [a], [b], [c] \in \mathbb{Z}_m$
- (c) $[1]$ ist neutrales Element der Multiplikation: $[1] \cdot [a] = [a] \quad \forall [a] \in \mathbb{Z}_m$

Beweis:

ähnlich elementar wie Beweis von Satz 7.8 auf Seite 64.

□

Bemerkung: Offensichtlich findet man nicht zu jedem $[a] \in \mathbb{Z}_m \setminus \{[0]\}$ ein inverses Element $[b]$ bzgl. der Multiplikation mit $[a] \cdot [b] = [1]$:
Beispiel 7.11 auf der vorherigen Seite zeigt, dass $[2], [3], [4]$ kein inverses Element haben.

7.13 Satz: (*Multiplikative inverse Elemente in $\mathbb{Z}_m$*)

$[a] \in \mathbb{Z}_m$ hat genau dann ein inverses Element bzgl. der Multiplikation, wenn a und m teilerfremd sind (d.h. $\text{ggT}(a, m) = 1$).

Beweis:

„ $\Rightarrow$ “ Sei $[b]$ ein multiplikatives Inverses zu $[a] \in \mathbb{Z}_m$

$$\begin{aligned} \Rightarrow [1] &= [a] \cdot [b] = [a \cdot b] \\ \Rightarrow \exists q \in \mathbb{Z} &: ab - 1 = qm \end{aligned} \quad (*)$$

Um zu zeigen, dass $\text{ggT}(a, m) = 1$, zeigen wir, dass für jeden Teiler c von a und m gilt: $c \mid 1$.

Für $c \mid a$ und $c \mid m$ folgt mit Satz 6.5 auf Seite 56(c):

$$c \mid ab - qm$$

Wegen (*) bedeutet dies: $c \mid 1$

„ $\Leftarrow$ “ siehe z.B. Beutelspacher / Zschiegner: Diskrete Mathematik für Einsteiger, 2002 (Satz 5.3.4)

□

Teil II

Algebra

Kapitel 8

Gruppen

8.1 Bedeutung in der Informatik

- Gruppen sind abstrakte Modelle für Mengen, auf denen eine Verknüpfung (etwa Addition oder Multiplikation) definiert ist.
- Allgemeine Aussagen aus der Gruppentheorie arbeiten die Gemeinsamkeiten hinter vielen Problemen heraus. (Beutelspacher: „Mathematik ist die Lehre von der guten Beschreibung“).

8.2 Definition: (Gruppe)

Eine Gruppe besteht aus einer Menge G und einer Verknüpfung $\cdot$, die je zwei Elementen aus G wieder ein Element aus G zuordnet, für die gilt:

- (a) Assoziativgesetz: $(a \cdot b) \cdot c = a \cdot (b \cdot c) \quad \forall a, b, c \in G.$
- (b) Neutrales Element: $\exists e \in G : e \cdot a = a \quad \forall a \in G.$
- (c) Inverse Elemente: Zu jedem $a \in G$ existiert $b \in G : b \cdot a = e.$
Man schreibt auch: $b =: a^{-1}$

Gilt zusätzlich

- (d) Kommutativgesetz: $a \cdot b = b \cdot a \quad \forall a, b \in G,$
so heißt $(G, \cdot)$ kommutative Gruppe (abelsche Gruppe).

$|G|$ heißt Ordnung der Gruppe. Ist $|G| < \infty$, spricht man von einer endlichen Gruppe.

Bemerkung: Erfüllt $(G, \cdot)$ lediglich das Assoziativgesetz, so spricht man von einer Halbgruppe (engl. semigroup). Halbgruppen mit neutralem Element heißen Monoide.

8.3 Beispiele

- (a) $(\mathbb{N}, +)$ ist eine Halbgruppe, aber kein Monoid, da kein neutrales Element existiert.
- (b) $(\mathbb{N}_0, +)$ ist ein Monoid (mit 0 als neutralem Element), jedoch keine Gruppe, da zu $a \in \mathbb{N}_0$ im Allgemeinen kein inverses Element existiert.
- (c) $(\mathbb{Z}, +), (\mathbb{Q}, +), (\mathbb{R}, +)$ sind kommutative Gruppen.
- (d) $(\mathbb{Z}, \cdot)$ ist ein Monoid (mit 1 als neutralem Element), aber keine Gruppe.
- (e) $(\mathbb{Q} \setminus \{0\}, \cdot), (\mathbb{R} \setminus \{0\}, \cdot)$ sind kommutative Gruppen.
- (f) Die Menge $\mathbb{Z}_m$ aller Restklassen modulo m bildet zusammen mit der modularen Addition eine kommutative Gruppe mit $[0]$ als neutralem Element (vgl. Satz 7.8 auf Seite 64).
- (g) Sei m eine Primzahl. Dann ist $(\mathbb{Z}_m \setminus \{[0]\}, \cdot)$ eine kommutative Gruppe mit neutralem Element $[1]$ (vgl. Satz 7.12 auf Seite 66 und Satz 7.13 auf Seite 66).
- (h) Die Menge aller Abbildungen $g : M \mapsto M$ bildet mit der Komposition ein Monoid (mit der identischen Abbildung als neutralem Element). Betrachtet man nur bijektive Abbildungen, liegt sogar eine (i.A. nicht-kommutative) Gruppe vor (vgl. § 5 auf Seite 45).
- (i) Spezialfall von (h):
Sei $M = \{1, \dots, n\}$. Dann bildet die Menge der bijektiven Abbildungen von $M \rightarrow M$ mit der Komposition die Permutationsgruppe (symmetrische Gruppe) $(S_n, \circ)$

Beispiel: für $n = 3$

$$\begin{aligned} \sigma_1 &:= \begin{pmatrix} 1 & 2 & 3 \\ 1 & 2 & 3 \end{pmatrix} & \sigma_2 &:= \begin{pmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \end{pmatrix} & \sigma_3 &:= \begin{pmatrix} 1 & 2 & 3 \\ 3 & 1 & 2 \end{pmatrix} \\ \sigma_4 &:= \begin{pmatrix} 1 & 2 & 3 \\ 2 & 1 & 3 \end{pmatrix} & \sigma_5 &:= \begin{pmatrix} 1 & 2 & 3 \\ 1 & 3 & 2 \end{pmatrix} & \sigma_6 &:= \begin{pmatrix} 1 & 2 & 3 \\ 3 & 2 & 1 \end{pmatrix} \end{aligned} \text{ Da-}$$

bei beschreibt z.B. σ_3 die Abbildung

$$\begin{aligned} 1 &\mapsto 3 \\ 2 &\mapsto 1 \\ 3 &\mapsto 2 \end{aligned}$$

$\sigma_3 \circ \sigma_4$ beschreibt Abbildung:

$$\begin{aligned} 1 &\xrightarrow{\sigma_4} 2 \xrightarrow{\sigma_3} 1 \\ 2 &\xrightarrow{\sigma_4} 1 \xrightarrow{\sigma_3} 3 \\ 3 &\xrightarrow{\sigma_4} 3 \xrightarrow{\sigma_3} 2 \end{aligned}$$

d.h. $\sigma_3 \circ \sigma_4 = \sigma_5$.

Die Verknüpfungstafel lautet

$\circ$	σ_1	σ_2	σ_3	σ_4	σ_5	σ_6
σ_1	σ_1	σ_2	σ_3	σ_4	σ_5	σ_6
σ_2	σ_2	σ_3	σ_1	σ_6	σ_4	σ_5
σ_3	σ_3	σ_1	σ_2	σ_5	σ_6	σ_4
σ_4	σ_4	σ_5	σ_6	σ_1	σ_2	σ_3
σ_5	σ_5	σ_6	σ_4	σ_3	σ_1	σ_2
σ_6	σ_6	σ_4	σ_5	σ_2	σ_3	σ_1

σ_1 ist neutrales Element.

Beachte: S_n ist nicht kommutativ: $\sigma_4 \cdot \sigma_6 \neq \sigma_6 \cdot \sigma_4$.

8.4 Satz: (Eindeutigkeit des neutr. Elements und der inversen Elemente)

In jeder Gruppe gibt es nur ein neutrales Element und jedes Element einer Gruppe hat genau ein inverses Element.

Beweis:

Sei e ein neutrales Element. Wir gehen in 4 Schritten vor:

- (a) Ist $b \cdot a = e$ (d.h. $b = a^{-1}$), so ist auch $a \cdot b = e$
(d.h. „linksinverse“ Elemente sind auch „rechtsinvers“).

Denn:

$$\begin{aligned}
 a \cdot b &= e \cdot (a \cdot b) \\
 &= (b^{-1} \cdot b)(a \cdot b) \\
 &= b^{-1}(\underbrace{b \cdot a}_e)b \quad (\text{Assoz.}) \\
 &= b^{-1}\underbrace{e \cdot b}_b \\
 &= b^{-1}b \\
 &= e
 \end{aligned}$$

- (b) Es gilt $ae = a \quad \forall a \in G$ (d.h. e ist auch „rechtsneutral“.)

Denn: $ae = a(a^{-1}a) \stackrel{(\text{Assoz.})}{=} (aa^{-1})a \stackrel{\text{nach (a)}}{=} ea = a$ (Def. des Linksneutralen)

- (c) Es gibt nur ein neutrales Element.

Denn: Sei e^* ein weiteres neutrales Element.

$$\Rightarrow e^*a = a \quad \forall a \in G \Rightarrow e = e^*e \stackrel{(b)}{=} e^* = e^*$$

- (d) Zu jedem $a \in G$ gibt es nur ein $a^{-1} \in G$.
Denn: Sei c ein weiteres inverses Element zu a .

$$\begin{aligned} \Rightarrow c \cdot a &= e & (*) \\ \Rightarrow c &\stackrel{\text{nach (b)}}{=} c \cdot e \\ &\stackrel{\text{nach (a)}}{=} c \cdot a \cdot a^{-1} \\ &\stackrel{\text{nach (*)}}{=} e \cdot a^{-1} \\ &= a^{-1} \end{aligned}$$

□

Gibt es Gruppen, die selbst wieder Gruppen enthalten ?

8.5 Satz: (Untergruppenkriterium)

Sei $(G, \cdot)$ eine (kommutative) Gruppe und $U \subset G$. Dann ist $(U, \cdot)$ genau dann eine (kommutative) Gruppe, wenn gilt:

$$a, b \in U \Rightarrow a \cdot b \in U \text{ und } a^{-1} \in U.$$

$(U, \cdot)$ heißt dann Untergruppe von $(G, \cdot)$.

Beweis:

Aus $a \cdot b \in U$ folgt, dass $\cdot$ abgeschlossen auf U abgeschlossen ist. Das Assoziativgesetz überträgt sich aus G . (ebenso das Kommutativgesetz im Fall einer kommutativen Gruppe)

Aus $a \cdot b \in U$ und $a^{-1} \in U$ folgt für alle $a \in U$:

$e = a^{-1}a \in U$ Ferner ist $a^{-1} \in U$ nach Voraussetzung.

□

8.6 Beispiel:

- (a) $(\mathbb{Z}, +)$ und $(\mathbb{Q}, +)$ sind Untergruppen von $(\mathbb{R}, +)$.
- (b) $(m\mathbb{Z}, +)$ mit $m\mathbb{Z} := \{mz \mid z \in \mathbb{Z}\}$ ist Untergruppe von $\mathbb{Z}$ für $m \in \mathbb{N}$.
- (c) Ist $(G, \cdot)$ eine Gruppe mit neutralem Element e , so sind $(\{e\}, \cdot)$ und $(G, \cdot)$ selbst wieder Untergruppen von $(G, \cdot)$.
- (d) Jede Permutation einer Menge $M = \{1, \dots, n\}$ lässt sich durch eine Sequenz von Vertauschungen zweier Elemente (Transpositionen) darstellen. Die Menge aller Permutationen von $M = \{1, \dots, n\}$ mit einer geraden Anzahl von Transpositionen bildet eine Untergruppe der symmetrischen

Gruppe $(S_n, \circ)$, die alternierende Gruppe $(A_n, \circ)$.
Z.B. besteht A_3 aus den Permutationen $\sigma_1, \sigma_2, \sigma_3$ (Notation aus 8.3 auf Seite 70(i)) und wird beschrieben durch die Verknüpfungstafel

$\circ$	σ_1	σ_2	σ_3
σ_1	σ_1	σ_2	σ_3
σ_2	σ_2	σ_3	σ_1
σ_3	σ_3	σ_1	σ_2

A_n ist sogar kommutativ (obwohl S_n nicht kommutativ ist).

8.7 Zyklenschreibweise für Permutationen

Sei σ eine Permutation vom $M = \{1, \dots, n\}$. Dann wird jedes Element von M nach höchstens n -maligen Anwenden von σ auf sich selbst abgebildet.

Beispiel:

$$\begin{aligned}
 \text{(a)} \quad \sigma_1 &:= \begin{pmatrix} 1 & 2 & 3 & 4 \\ 3 & 4 & 2 & 1 \end{pmatrix} & 1 \xrightarrow{\sigma_1} 3 \xrightarrow{\sigma_1} 2 \xrightarrow{\sigma_1} 4 \xrightarrow{\sigma_1} 1 \\
 \text{(b)} \quad \sigma_2 &:= \begin{pmatrix} 1 & 2 & 3 & 4 \\ 3 & 1 & 2 & 4 \end{pmatrix} & 1 \xrightarrow{\sigma_2} 3 \xrightarrow{\sigma_2} 2 \xrightarrow{\sigma_2} 1 \\
 & & 4 \xrightarrow{\sigma_2} 4 \\
 \text{(c)} \quad \sigma_3 &:= \begin{pmatrix} 1 & 2 & 3 & 4 \\ 3 & 4 & 1 & 2 \end{pmatrix} & 1 \xrightarrow{\sigma_3} 3 \xrightarrow{\sigma_3} 1 \\
 & & 2 \xrightarrow{\sigma_3} 4 \xrightarrow{\sigma_3} 2
 \end{aligned}$$

Das motiviert die Zyklenschreibweise

- (a) $\sigma_1 = (1324)$
- (b) $\sigma_2 = (132)$ (Der Einerzyklus (4) wird weggelassen.)
- (c) $\sigma_3 = (13)(24)$

8.8 Definition: (Nebenklassen, Rechts- und Linksnebenklassen)

Mit Hilfe von Äquivalenzrelationen kann man eine Menge in Äquivalenzklassen zerlegen.

(Bsp.: $\mathbb{Z}$ wird in m Äquivalenzklassen modulo m zerlegt.)

Gibt es eine ähnliche Zerlegung bei Gruppen?

Def.: Sei $(G, \cdot)$ eine Gruppe mit Untergruppe $(U, \cdot)$. Ferner sei $g \in G$. Dann nennen wir

$$\begin{aligned}
 gU &:= \{g \cdot u \mid u \in U\} \text{ Linksnebenklasse von } g \text{ bzgl. } U, \\
 Ug &:= \{u \cdot g \mid u \in U\} \text{ Rechtsnebenklasse von } g \text{ bzgl. } U.
 \end{aligned}$$

Bemerkung: Häufig betrachtet man nur Linksnebenklassen und nennt diese dann Nebenklassen.

8.9 Satz: (Nebenklassenzerlegung einer Gruppe)

Sei $(G, \cdot)$ eine Gruppe, $g, h \in G$ und $(U, \cdot)$ eine Untergruppe. Dann gilt: für $g, h \in G$:

- (a) $g \in U \Rightarrow gU = U$
- (b) Zwei (Links-)Nebenklassen gU, hU sind entweder gleich oder disjunkt.
- (c) Jedes $a \in G$ liegt in einer eindeutig bestimmten (Links-)Nebenklasse, d.h. die Nebenklassen von U bilden eine Partition von G .
- (d) Alle (Links-)Nebenklassen bzgl. einer festen Untergruppe U sind gleichmächtig:
 $|gU| = |U| \quad \forall g \in G$.

Beweis:

- (a) Aus der Abgeschlossenheit folgt $gU \subset U$. Mit (d) folgt schliesslich $gU = U$.
- (b) Ann.: gU und hU haben ein gemeinsames Element
 $\Rightarrow \exists a, b \in U : ga = hb. \quad (*)$
 $\Rightarrow gU \stackrel{(a)}{=} g(aU) = (ga)U \stackrel{(*)}{=} (hb)U = h(bU) \stackrel{(a)}{=} hU$
- (c) $a \in G$ liegt in aU , da U das neutrale Element e enthält. Die Eindeutigkeit folgt aus (b).
- (d) Es gilt: $|gU| = |U|$, denn:
 - Zu jedem $a \in U$ ex. Element $ga \in gU$, d.h. $|U| \leq |gU|$.
 - Zu jedem $ga \in gU$ existiert $a \in U$, d.h. $|gU| \leq |U|$.

□

8.10 Beispiele

- (a) $(5\mathbb{Z}, +)$ ist eine Untergruppe von $(\mathbb{Z}, +)$. Wir können $\mathbb{Z}$ in 5 Nebenklassen zerlegen.

$$\begin{aligned} 0 + 5\mathbb{Z} &= [0] \\ 1 + 5\mathbb{Z} &= [1] \\ 2 + 5\mathbb{Z} &= [2] \\ 3 + 5\mathbb{Z} &= [3] \\ 4 + 5\mathbb{Z} &= [4] \end{aligned}$$

Dies sind gerade die 5 Kongruenzklassen modulo 5.

- (b) Die alternierende Gruppe $(A_3, \circ)$ ist eine Untergruppe der symmetrischen Gruppe $(S_3, \circ)$ (Vgl. 8.3 auf Seite 70(i), 8.6 auf Seite 72(d)).
 $S_3 := \{\sigma_1, \sigma_2, \dots, \sigma_6\}$ hat die Linksnebenklassen:

$$\begin{aligned} A_3 &= \{\sigma_1, \sigma_2, \sigma_3\} & (= \sigma_1 A_3 = \sigma_2 A_3 = \sigma_3 A_3) \\ \sigma_4 A_3 &= \{\sigma_4, \sigma_5, \sigma_6\} & (= \sigma_5 A_3 = \sigma_6 A_3) \end{aligned}$$

$$\begin{array}{lll} \text{denn: } \sigma_4 \sigma_1 = \sigma_4 & \sigma_5 \sigma_1 = \sigma_5 & \sigma_6 \sigma_1 = \sigma_6 \\ \sigma_4 \sigma_2 = \sigma_5 & \sigma_5 \sigma_2 = \sigma_6 & \sigma_6 \sigma_2 = \sigma_4 \\ \sigma_4 \sigma_3 = \sigma_6 & \sigma_5 \sigma_3 = \sigma_4 & \sigma_6 \sigma_3 = \sigma_5 \end{array}$$

Wie viele Nebenklassen besitzt eine Nebenklassenzerlegung von G bzgl. einer Untergruppe U ?

8.11 G modulo U

Sei $(U, \cdot)$ eine Untergruppe von $(G, \cdot)$. Dann bezeichnen wir die Menge aller Linksnebenklassen mit G/U (gesprochen: „ G modulo U “), und $G : U := |G/U|$ heißt Index von U in G .

8.12 Satz: (von Lagrange)

Sei $(G, \cdot)$ eine endliche Gruppe mit Untergruppe $(U, \cdot)$. Dann ist die Untergruppenordnung $|U|$ Teiler der Gruppenordnung $|G|$. Für die Zahl der Linksnebenklassen gilt:

$$G : U = \frac{|G|}{|U|}$$

Beweis:

Nach Satz 8.9 auf der vorherigen Seite sind alle Nebenklassen von G bzgl. U gleichmächtig und bilden eine Partition von G . Also muss $|U|$ Teiler von $|G|$ sein und $\frac{|G|}{|U|}$ ist die Zahl der Nebenklassen.

□

8.13 Beispiel:

- (a) $|S_3| = 6, |A_3| = 3$
 $S_3/A_3 = \{A_3, \sigma_4 A_3\}$
 $S_3 : A_3 = \frac{|S_3|}{|A_3|} = \frac{6}{3} = 2$

- (b) Eine Gruppe mit 30 Elementen kann nur Untergruppen mit 1, 2, 3, 5, 6, 10, 15, 30 Elementen besitzen.

8.14

Wichtig sind Untergruppen $(N, \cdot)$ einer Gruppe $(G, \cdot)$ für die Links- und Rechtsnebenklassen identisch sind:

$$gN = Ng \quad \forall g \in G$$

Sie heißen Normalteiler. Offensichtlich ist in einer kommutativen Gruppe $(G, \cdot)$ jede Untergruppe ein Normalteiler.

Warum sind Normalteiler wichtig?

Ist eine Untergruppe $(N, \cdot)$ Normalteiler von $(G, \cdot)$, so kann man zeigen, dass die Nebenklassenmenge G/N zu einer Gruppe wird, indem man sie mit der Verknüpfung:

$$(gN)(hN) := (gh)N$$

ausstattet. Diese Gruppe heißt Faktorgruppe von G nach N .

8.15 Beispiel:

Die Untergruppe $(6\mathbb{Z}, +)$ ist Normalteiler in $(\mathbb{Z}, +)$, da $(\mathbb{Z}, +)$ eine kommutative Gruppe ist. In 7.6 auf Seite 63 hatten wir auf der Restklassenmenge $\mathbb{Z}_6 = \mathbb{Z}/6\mathbb{Z}$ die modulare Addition definiert durch:

$$[a] + [b] := [a + b].$$

Wegen $[a] = a + 6\mathbb{Z}$, $[b] = b + 6\mathbb{Z}$ ist dies nichts anderes als die in 8.14 eingeführte Verknüpfung

$$(a + 6\mathbb{Z}) + (b + 6\mathbb{Z}) := (a + b) + 6\mathbb{Z}.$$

8.16 Definition: (Abbildung zwischen Gruppen)

Seien $(G_1, \circ)$, $(G_2, \bullet)$ Gruppen.

- (a) Ein Homomorphismus von G_1 nach G_2 ist eine Abbildung

$$f : G_1 \rightarrow G_2 \text{ mit } \underbrace{f(a \circ b)}_{\text{Verknüpfung in } G_1} = \underbrace{f(a) \bullet f(b)}_{\text{Verknüpfung in } G_2} \quad \forall a, b \in G$$

- (b) Ein injektiver Homomorphismus heißt Monomorphismus.

- (c) Ein surjektiver Homomorphismus heißt Epimorphismus.
- (d) Ein bijektiver Homomorphismus heißt Isomorphismus.
Man schreibt: $G_1 \simeq G_2$ („ G_1 isomorph G_2 “).
- (e) Ein Homomorphismus von G_1 in sich selbst heißt Endomorphismus.
- (f) Ein Isomorphismus von G_1 in G_1 heißt Automorphismus.

8.17 Definition: (Bild, Kern)

Sei $f : G_1 \rightarrow G_2$ ein Homomorphismus der Gruppen G_1, G_2 . Dann heißt

$$\text{Im}(f) := \{f(g_1) \mid g_1 \in G_1\}$$

das Bild von f . Sei ferner e_2 das neutrale Element von $(G_2, \cdot)$. Dann nennt man

$$\text{Ker}(f) := \{g_1 \in G_1 \mid f(g_1) = e_2\}$$

den Kern von f .

8.18 Satz: (Homomorphismus für Gruppen)

Warum ist der Kern eines Homomorphismus wichtig ? Man kann zeigen:

Sei $f : G_1 \rightarrow G_2$ ein Homomorphismus der Gruppen G_1 und G_2 . Dann ist $\text{Ker}(f)$ Normalteiler von f und die Faktorgruppe $G_1/\text{Ker}(f)$ ist isomorph zum Bild von f :

$$G_1/\text{Ker}(f) \simeq \text{Im}(f).$$

Bemerkung: Man kann also eine nicht bijektive Abbildung bijektiv machen, indem man zum Faktorraum übergeht, d.h. Elemente ignoriert, die auf das neutrale Element von G_2 abgebildet werden.

Kapitel 9

Ringe

9.1 Motivation

- Häufig gibt es auf einer Menge 2 Verknüpfungen: eine „Addition“ und eine „Multiplikation“.
(**Beispiel:** $(\mathbb{Z}, +, \cdot)$)
- In $(\mathbb{Z}, +, \cdot)$ gibt es sogar noch eine Division mit Rest.
- Lassen sich diese Konzepte algebraisch abstrahieren?

9.2 Definition: *(Ring)*

Eine Menge R mit zwei Verknüpfungen $+, \cdot$ auf R heißt Ring, wenn gilt:

- (a) $(R, +)$ ist eine kommutative Gruppe.
- (b) $(R, \cdot)$ ist eine Halbgruppe.
- (c) *Distributivgesetze:*

$$\left. \begin{array}{l} a \cdot (b + c) = (a \cdot b) + (a \cdot c) \\ (b + c) \cdot a = (b \cdot a) + (c \cdot a) \end{array} \right\} \forall a, b, c \in \mathbb{R}.$$

Ist $(R, \cdot)$ sogar ein Monoid, so heißt $(R, +, \cdot)$ Ring mit Einselement. Gilt neben (a)-(c) zusätzlich:

- (d) *Kommutativgesetz d. Multiplikation:*

$$a \cdot b = b \cdot a \quad \forall a, b \in \mathbb{R}$$

so heißt $(R, +, \cdot)$ kommutativer Ring.

9.3 Konventionen

In einem Ring $(R, +, \cdot)$ nennt man oft

- das neutrale Element der Addition Nullelement (0) .
- das neutrale Element der Multiplikation Einselement (1) .
- das additive Inverse zu $a : -a$.
- das multiplikative Inverse zu $a : \frac{1}{a}$.

Um Klammern zu sparen, vereinbart man „Punkt vor Strich“.

9.4 Beispiele

- (a) $(\mathbb{Z}, +, \cdot)$ ist ein kommutativer Ring mit Eins (ebenso wie $(\mathbb{Q}, +, \cdot)$ und $(\mathbb{R}, +, \cdot)$).
- (b) Die Menge $\mathbb{Z}_m = \mathbb{Z}/m\mathbb{Z}$ der Restklassen modulo m bildet zusammen mit der modularen Addition und Multiplikation den Restklassenring $(\mathbb{Z}/m\mathbb{Z}, +, \cdot)$. Auch er ist kommutativ. Er besitzt das Einselement $[1]$ (vgl. § 7 auf Seite 61).

Nachweis des ersten Distributivgesetzes:

$$\begin{aligned}[a] \cdot ([b] + [c]) &= [a] \cdot [b + c] \\ &= [a \cdot (b + c)] \\ &= [a \cdot b + a \cdot c] \\ &= [a \cdot b] + [a \cdot c] \\ &= [a] \cdot [b] + [a] \cdot [c] \\ &\quad \forall [a], [b], [c] \in \mathbb{Z}/m\mathbb{Z}\end{aligned}$$

Zweites Distributivgesetz folgt aus Kommutativität.

Beispiel für nicht kommutative Ringe folgt später (Matrizen).

9.5 Satz: (Unterringkriterium)

Sei $(R, +, \cdot)$ ein Ring und $S \subset R$. Dann ist $(S, +, \cdot)$ genau dann ein Ring, falls

- (a) $(S, +)$ ist Untergruppe von $(R, +)$.

(b) $(S, \cdot)$ ist abgeschlossen: $a, b \in S \Rightarrow a \cdot b \in S$.

$(S, +, \cdot)$ heißt dann Unterring von $(R, +, \cdot)$.

Beweis:

Wie Satz 8.5.

9.6 Beispiel:

$(m\mathbb{Z}, +, \cdot)$ ist Unterring in $(\mathbb{Z}, +, \cdot)$, denn

(a) $(m\mathbb{Z}, +)$ ist Untergruppe von $(\mathbb{Z}, +)$.

(b) $(m\mathbb{Z}, \cdot)$ ist abgeschlossen:

Seien $a, b \in m\mathbb{Z}$:

$\exists q_1, q_2 \in \mathbb{Z}$:

$a = q_1 m, \quad b = q_2 m$

$\Rightarrow a \cdot b = (q_1 m)(q_2 m) = \underbrace{(q_1 q_2 m)}_{\in \mathbb{Z}} \quad m \in m\mathbb{Z}.$

9.7 Polynomringe

Sie sind die wichtigsten Ringe. Sei $(R, +, \cdot)$ ein Ring (z.B. $R = \mathbb{R}$) und $a_0, a_1, \dots, a_n \in \mathbb{R}$. Dann nennen wir die Abbildung:

$$p: R \rightarrow R, \quad x \mapsto \sum_{k=0}^n a_k x^k$$

Polynom (über R). Dabei ist $x^k := \underbrace{x \cdot x \cdot \dots \cdot x}_{k\text{-mal}}$. $a_0, \dots, a_n$ heißen Koeffizienten

von p . Ist $a_0 \neq 0$, so heißt n Grad von $p: n = \deg(p)$.

(**Beispiel:** $p(x) = 5x^3 - 1,3x + 6$ ist Polynom 3. Grades.)

Die Menge aller Polynom über R nennen wir $R[x]$.

Auf $R[x]$ definieren wir eine Addition und eine Multiplikation „punktweise“ durch

$$\begin{aligned} (p + q)(x) &:= p(x) + q(x) \\ (p \cdot q)(x) &:= p(x) \cdot q(x). \end{aligned}$$

Dann ist $(R[x], +, \cdot)$ ein Ring, der Polynomring über R .

Der Nachweis der Ringeigenschaften ist arbeitsaufwändig, aber nicht schwierig. Man verwendet:

Mit $p(x) = \sum_{k=0}^n a_k x^k$, $q(x) = \sum_{k=0}^n b_k x^k$ gilt:

$$(p+q)(x) = \sum_{k=0}^n (a_k + b_k) x^k$$
$$(p \cdot q)(x) = \sum_{k=0}^{n+n} \left(\sum_{\substack{i+j=k \\ 0 \leq i, j \leq n}} a_i b_j \right) x^k.$$

Dabei wurde $n := \max(\deg(p), \deg(q))$ gewählt und entsprechend Koeffizienten des Polynoms kleineren Grades Null gesetzt.

9.8 Horner-Schema

Oft müssen Funktionswerte von Polynomen effizient berechnet werden: z.B. approximiert der Taschenrechner die Sinusfunktion durch ein Polynom. Eine naive Auswertung eines Polynoms

$$p(x) = 5x^7 + 4x^6 - 3x^5 + 4x^4 + 6x^3 - 7x^2 + 4x - 1$$

erfordert 7 Additionen und $7+6+5+4+3+2+1 = 28$ Multiplikationen.

Wesentlich effizienter ist das Horner-Schema, das eine geschickte Klammerung ausnutzt:

$$p(x) = ((((((5x + 4)x - 3)x + 4)x + 6)x - 7)x + 4)x - 1$$

Arbeitet man die Klammern von innen nach außen ab, benötigt man 7 Additionen und 7 Multiplikationen.

Allgemein benötigt man bei der naiven Auswertung eines Polynoms n -ten Grades n Additionen und $\sum_{k=1}^n k = \frac{n(n+1)}{2}$ Multiplikationen. Das Horner-Schema benötigt nur n Multiplikationen.

9.9 Polynomdivision

Im Ring $(\mathbb{Z}, +, \cdot)$ haben wir in 6 auf Seite 55 Division mit Rest betrachtet. Kann man Ähnliches auch im Polynomring $(\mathbb{R}[x], +, \cdot)$ tun? (**Beachte:** $R = \mathbb{R}$ hier!!!)

Satz:

Polynomdivision In $(\mathbb{R}[x], +, \cdot)$ kann man eine Division mit Rest durchführen: Zu $a, b \in \mathbb{R}[x]$ mit $b \neq 0$ existieren eindeutig bestimmte $q, r \in \mathbb{R}[x]$ mit $a = qb + r$ und $\deg(r) < \deg(b)$.

9.10 Praktische Durchführung der Polynomdivision

Analog zur Division natürlicher Zahlen:

$$\begin{array}{r} 365 : 7 = 52 \text{ Rest } 1 \\ -35 \\ \hline 15 \\ -14 \\ \hline 1 \end{array}$$

führt man die Polynomdivision durch:

$$\begin{array}{r} (x^4 + 2x^3 + 3x^2 + 4x + 5) : (x^2 + 1) = x^2 + 2x + 2 + \frac{2x+3}{x^2+1} \\ -x^4 - x^2 \\ \hline 2x^3 + 2x^2 + 4x \\ -2x^3 - 2x \\ \hline 2x^2 + 2x + 5 \\ -2x^2 - 2 \\ \hline 2x + 3 \end{array}$$

bzw.: $(x^4 + 2x^3 + 3x^2 + 4x + 5) : (x^2 + 1) = x^2 + 2x + 2 \text{ Rest } 2x + 3$.

Satz 9.9 auf der vorherigen Seite hat eine wichtige Folgerung.

9.11 Satz: (Abspaltung von Nullstellen)

Hat $p \in \mathbb{R}[x]$ die Nullstelle x_0 (d.h. $p(x_0) = 0$), so ist p durch das Polynom $(x - x_0)$ ohne Rest teilbar.

Beweis:

Nach Satz 9.9 auf der vorherigen Seite existiert für $p(x)$ und $b(x) := x - x_0$ eine Division mit Rest:

$$\exists q(x), r(x) \text{ mit } p(x) = q(x)b(x) + r(x).$$

$$\text{In } x_0 \text{ gilt dann } 0 = p(x_0) = q(x_0) \underbrace{b(x_0)}_0 + r(x_0) \Rightarrow r(x_0) = 0$$

Wegen $1 = \deg(b) > \deg(r)$ ist $\deg(r) = 0 \Rightarrow r = a_0 = 0$.

□

9.12 Satz: (Zahl der Nullstellen)

Ein von Null verschiedenes Polynom $p \in \mathbb{R}[x]$ vom Grad n hat höchstens n Nullstellen.

Beweis:

Annahme: $p(x)$ hat mehr als n Nullstellen. Sukzessives Abspalten der n Nullstellen $x_1, x_2, \dots, x_n$ ergibt:

$\exists q \in \mathbb{R}[x] : p(x) = (x - x_1) \cdot (x - x_2) \cdot \dots \cdot (x - x_n) \cdot q(x)$. $q(x)$ hat Grad 0, sonst wäre $\deg(p) > n$.

D.h. $q(x) = a_0$. Sei x_{n+1} eine weitere Nullstelle.

$$0 = p(x_{n+1}) = (x_{n+1} - x_1) \cdot \dots \cdot (x_{n+1} - x_n) \cdot q(x_{n+1})$$

Also ist einer der Faktoren $(x_{n+1} - x_1), \dots, (x_{n+1} - x_n), \underbrace{q(x_{n+1})}_{=a_0}$ Null. $a_0 = 0$

geht nicht, da $p(x) \neq 0$ vorausgesetzt.

$\Rightarrow x_{n+1}$ stimmt mit einem $x_n (k = 1, \dots, n)$ überein.

□

Im Ring $(\mathbb{Z}, +, \cdot)$ haben wir in 6 auf Seite 55 Teilbarkeit und ggT studiert. Geht dies auch im Polynomring $(\mathbb{R}[x], +, \cdot)$?

9.13 Definition: (ggT von Polynomen)

Seien $a, b \in \mathbb{R}[x]$. Dann ist a durch b teilbar, wenn es ein $q \in \mathbb{R}[x]$ gibt mit $a(x) = q(x)b(x)$. Ein Polynom $p \in \mathbb{R}[x]$ ist gemeinsamer Teiler von a und b , falls p sowohl a als auch b teilt. p heißt größter gemeinsamer Teiler von a und b (ggT(a, b)), falls p durch jeden gemeinsamen Teiler von a und b teilbar ist.

9.14 Euklidischer Algorithmus für Polynome

Analog zu 6.11 auf Seite 59 bestimmt man den ggT zweier Polynome $a, b \in \mathbb{R}[x]$ mit $\deg(a) \geq \deg(b)$ mit dem Euklidischen Algorithmus.

Beispiel:

$$\begin{aligned} a(x) &= x^4 + x^3 - x^2 + x + 2 \\ b(x) &= x^3 + 2x^2 + 2x + 1 \end{aligned}$$

Polynomdivision ergibt:

$$\begin{aligned}(x^4 + x^3 - x^2 + x + 2) &= (x - 1)(x^3 + 2x^2 + 2x + 1) + (-x^2 + 2x + 3) \\(x^3 + 2x^2 + 2x + 1) &= (-x - 4)(-x^2 + 2x + 3) + (13x + 13) \\(-x^2 + 2x + 3) &= \left(-\frac{1}{13}x + \frac{3}{13}\right)(13x + 13)\end{aligned}$$

Also ist $\text{ggT}(a(x), b(x)) = (13x + 13)$.

Ebenso wie jedes Polynom $c \cdot (13x + 13)$ mit einer Konstante $c \neq 0$.

Gibt es „Primfaktoren“ in Polynomringen?

9.15 Definition: (*reduzibel, irreduzibel*)

Sei $(R, +, \cdot)$ ein kommutativer Ring mit 1. Ein Polynom $p \in R[x]$ heißt reduzibel über R , wenn es nichtkonstante Polynome $a, b \in R[x]$ gibt mit $p = a(x) \cdot b(x)$. Andernfalls heißt p irreduzibel über R .

9.16 Beispiele

- (a) $x^2 - 4$ ist reduzibel über $\mathbb{Z}$: $x^2 - 4 = (x + 2)(x - 2)$
- (b) $x^2 - 3$ ist irreduzibel über $\mathbb{Z}$ und $\mathbb{Q}$, aber reduzibel über $\mathbb{R}$: $x^2 - 3 = (x + \sqrt{3})(x - \sqrt{3})$
- (c) $x^2 + 1$ ist irreduzibel über $\mathbb{R}$

9.17 Irreduzible Polynome als „Primfaktoren“ in $(\mathbb{R}[x], +, \cdot)$

Irreduzible Polynome in $(\mathbb{R}[x], +, \cdot)$ haben ähnliche Eigenschaften wie Primfaktoren in $(\mathbb{Z}, +, \cdot)$. So existiert z.B. eine „Primfaktorzerlegung“:

Jedes nichtkonstante Polynom $p \in \mathbb{R}[x]$ lässt sich als Produkt irreduzibler Polynome aus $\mathbb{R}[x]$ darstellen. Die Darstellung ist bis auf die Reihenfolge der irreduziblen Polynome und bis auf Multiplikation mit konstanten Polynomen eindeutig. (vgl. Satz 6.6 auf Seite 56)

9.18 Beispiel:

$p(x) = x^3 - x^2 + x - 1$ hat in $\mathbb{R}[x]$ die folgende Zerlegung in irreduzible Polynome:

$$p(x) = (x^2 + 1)(x - 1)$$

Äquivalent sind z.B.

$$p(x) = (x^2 + 1)(x - 1) = \left(\frac{x^2}{4} + \frac{1}{4}\right)(4x - 4)$$

Generell kann man zeigen, dass es über $\mathbb{R}$ nur 2 Typen von irreduziblen Polynomen gibt:

- (a) lineare Polynome (d.h. Grad 1).
- (b) quadratische Polynome $ax^2 + bx + c$ mit $b^2 - 4ac < 0$.

Kapitel 10

Körper

10.1 Definition: (Körper)

Ein Körper $(K, +, \cdot)$ besteht aus einer Menge K und zwei Verknüpfungen $+$ und $\cdot$ auf K , für die gilt:

- (a) $(K, +, \cdot)$ ist ein kommutativer Ring mit 1.
- (b) Inverse Elemente: Zu jedem $a \in K$ mit $a \neq 0$ existiert ein $a^{-1} \in K$ mit $a^{-1} \cdot a = 1$.

10.2 Bemerkung:

- (a) $(K, +, \cdot)$ besteht somit aus den kommutativen Gruppen $(K, +)$ und $(K \setminus \{0\}, \cdot)$, zwischen denen ein Distributivgesetz gilt.
- (b) Statt $K \setminus \{0\}$ schreibt man oft K^* .
- (c) Im Englischen heißen Körper fields.
- (d) Falls $(K^*, \cdot)$ nur eine nicht kommutative Gruppe ist, bezeichnet man $(K, +, \cdot)$ als Schiefkörper.

10.3 Beispiele

- (a) $(\mathbb{Z}, +, \cdot)$ ist kein Körper, da zu $a \in \mathbb{Z}^*$ i.A. kein a^{-1} existiert (bzgl. $\cdot$).
- (b) $(\mathbb{Q}, +, \cdot)$ und $(\mathbb{R}, +, \cdot)$ sind Körper.

Weitere Beispiele folgen später.

Welche Eigenschaften gelten in Körpern?

10.4 Satz: (Eigenschaften von Körpern)

In einem Körper $(K, +, \cdot)$ gilt:

(a) $a \cdot 0 = 0 \quad \forall a \in K$

(b) Nullteilerfreiheit:
Sind $a, b \in K$ mit $a \neq 0$ und $b \neq 0$, folgt $a \cdot b \neq 0$ (d.h. aus $a \cdot b = 0$ folgt $a = 0$ oder $b = 0$).

Beweis:

(a) Da 0 neutrales Element bzgl. $+$ ist, gilt:

$$a \cdot 0 + 0 = a \cdot 0 = a \cdot (0 + 0) \stackrel{\text{Distr.}}{=} a \cdot 0 + a \cdot 0$$

Addition von $-a \cdot 0$ auf beiden Seiten ergibt die Behauptung
 $0 = a \cdot 0$.

(b) Seien $a \neq 0, b \neq 0$.

Annahme:

$$\begin{aligned} a \cdot b &= 0 \\ \Rightarrow b &= (a^{-1}a) \cdot b \\ &= a^{-1} \cdot \underbrace{(ab)}_0 \\ &= a^{-1} \cdot 0 \stackrel{(a)}{=} 0 \end{aligned}$$

Widerspruch zu $b \neq 0$

□

Bemerkung: 10.4 (a) impliziert, dass man in Körpern nicht durch 0 dividieren darf: $q := \frac{a}{b}$ bedeutet $a = q \cdot b$. Ist $b = 0$, folgt $a = 0$. Für $a = 0$ ist jedoch $a = q \cdot b$ für jedes $q \in K$ erfüllt. Somit macht auch $\frac{0}{0}$ keinen Sinn.

10.5 Endliche Körper

Wir betrachten den Restklassenring $\mathbb{Z}_m = \mathbb{Z}/m\mathbb{Z}$ mit der modularen Addition und Multiplikation. In 8.3 auf Seite 70(g) haben wir gesehen (vgl. auch 7.12 auf Seite 66 und 7.13 auf Seite 66):

$(\mathbb{Z}_m^*, \cdot)$ ist eine kommutative Gruppe, falls m eine Primzahl ist.

Man kann sogar zeigen:

$(\mathbb{Z}_m, +, \cdot)$ ist genau dann ein Körper, wenn m Primzahl ist.

Beispiel 7.11 illustriert, dass $(\mathbb{Z}_6, +, \cdot)$ kein Körper sein kann, da z.B. $[2]$ kein multiplikatives Inverses hat und $\mathbb{Z}_6$ nicht nullteilerfrei ist.

Allgemein bezeichnet man einen endlichen Körper mit q Elementen als Galoisfeld $\text{GF}(q)$ (Evariste Galois, 1811 - 1832). Galoisfelder existieren nur für Primzahlpotenzen $q = p^m$ und sind nur für festes q eindeutig (bis auf Isomorphie). Für eine Primzahl p gilt also:

$$\text{GF}(p) = \mathbb{Z}_p$$

10.6 Der Körper der komplexen Zahlen

Problem:

- Manche Polynome haben in $\mathbb{R}$ keine Nullstellen.
- **Beispiel:** $p(x) = x^2 + 1$.
- Kann man $\sqrt{-1}$ sinnvoll definieren?

Abhilfe:

- Wir betten $\mathbb{R}$ in $\mathbb{C} := \mathbb{R}^2 = \{(a, b) \mid a, b \in \mathbb{R}\}$ ein, indem wir $\mathbb{R}$ als x -Achse $\{(a, 0) \mid a \in \mathbb{R}\}$ interpretieren. Dort muss gelten:
 - Addition: $(a, 0) + (c, 0) = (a + c, 0)$
 - Multiplikation: $(a, 0) \cdot (c, 0) = (a \cdot c, 0) \quad \forall a, b \in \mathbb{R}$
- Wir suchen eine Erweiterung der Addition und Multiplikation auf $\mathbb{C}$, so dass
 - $(\mathbb{C}, +, \cdot)$ ein Körper ist und
 - eine Zahl $(a, b) \in \mathbb{C}$ existiert mit $(a, b) \cdot (a, b) = (-1, 0)$.

10.7 Definition: (Komplexe Zahlen)

Auf $\mathbb{C} := \mathbb{R}^2 = \{(a, b) \mid a, b \in \mathbb{R}\}$ sind folgende Verknüpfungen definiert:

- (a) Addition: $(a, b) + (c, d) := (a + c, b + d) \quad \forall a, b, c, d \in \mathbb{R}.$
- (b) Multiplikation: $(a, b) \cdot (c, d) := (a \cdot c - b \cdot d, a \cdot d + b \cdot c) \quad \forall a, b, c, d \in \mathbb{R}.$

10.8 Konsequenzen aus Def. 10.7

- Einschränkung auf x -Achse liefert Addition/Multiplikation der reellen Zahlen:

$$\begin{aligned}(a, 0) + (c, 0) &= (a + c, 0 + 0) = (a + c, 0) \\ (a, 0) \cdot (c, 0) &= (a \cdot c - 0 + 0, a \cdot 0 + 0 \cdot c) = (a \cdot c, 0) \quad \forall a, c \in \mathbb{R}.\end{aligned}$$

- Das neutrale Element der Multiplikation auf $\mathbb{C}$ ist $(1, 0)$:

$$(1, 0) \cdot (c, d) = (c, d).$$

- In $\mathbb{C}$ existiert $\sqrt{-1}$, denn:

$$(0, 1) \cdot (0, 1) = (0 \cdot 0 - 1 \cdot 1, 0 \cdot 1 + 1 \cdot 0) = (-1, 0).$$

Für $(0, 1)$ schreiben wir **i** (imaginäre Einheit).

10.9 Satz: (Körper der komplexen Zahlen)

$(\mathbb{C}, +, \cdot)$ bildet einen Körper, den Körper der komplexen Zahlen.

Beweis:

Elementar, abgesehen von der Existenz des inversen Elements der Multiplikation.

Für $(a, b) \neq (0, 0)$ gilt:

$$\begin{aligned}(a, b^{-1}) &:= \left(\frac{a}{a^2 + b^2}, \frac{-b}{a^2 + b^2} \right) \\ \left(\frac{a}{a^2 + b^2}, \frac{-b}{a^2 + b^2} \right) (a, b) &= \left(\frac{a^2}{a^2 + b^2} + \frac{b^2}{a^2 + b^2}, \frac{ab}{a^2 + b^2} - \frac{ab}{a^2 + b^2} \right) \\ &= (1, 0)\end{aligned}$$

□

10.10 Praktisches Rechnen mit komplexen Zahlen

Statt (a, b) schreibe man $a + i \cdot b$, verwendet $i^2 = -1$ und rechnet ansonsten wie mit reellen Zahlen. Damit erhält man direkt Def. 10.7 auf der vorherigen Seite der Addition/Multiplikation:

Addition:

$$(a + ib) + (c + id) = (a + c) + i(b + d)$$

Multiplikation

$$\begin{aligned}(a + ib) \cdot (c + id) &= ac + iad + ibc + \underbrace{i^2}_{-1} bd \\ &= (ac - bd) + i(ad + bc)\end{aligned}$$

10.11 Definition: (komplex konjugiertes Element, Betrag, Realteil, Imaginärteil)

Zu einer komplexen Zahl $z = a + ib$ definiert man $\bar{z} := a - ib$ als das (komplex) konjugierte Element zu z .

Ferner nennt man $|z| = \sqrt{z \cdot \bar{z}} = \sqrt{a^2 + b^2}$ (siehe auch 10.12) den Betrag von z .

a heißt Realteil, b heißt Imaginärteil von z .

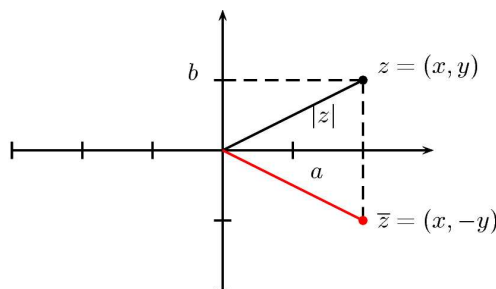
10.12 Geometrische Interpretation

Betrachtet man z als Vektor $(a, b) \in \mathbb{R}^2$, so ist $\bar{z}$ der an der x -Achse gespiegelte Vektor.

Wegen

$$\begin{aligned}|z| &= \sqrt{(z \cdot \bar{z})} \\ &= \sqrt{(a + ib)(a - ib)} \\ &= \sqrt{(a^2 - i^2 b^2)} \\ &= \sqrt{a^2 + b^2}\end{aligned}$$

ist $|z|$ die Länge des Vektors (a, b) .



10.13 Nützlichkeit des konjugiert komplexen Elementes

Z.B. um bei komplexen Brüchen $\frac{a+ib}{c+id}$ den Nenner reell zu machen. Man erweitert mit $c-id$:

$$\begin{aligned}\frac{a+ib}{c+id} \cdot \frac{c-id}{c-id} &= \frac{ac-iad+ibc-i^2bd}{c^2-cid+idc-i^2d^2} \\ &= \frac{ac+bd}{c^2+d^2} + i^2 \frac{bc-ad}{c^2+d^2}.\end{aligned}$$

Welche Bedeutung hat $\mathbb{C}$ für die Nullstellen von Polynomen? Man kann zeigen:

10.14 Satz: (Fundamentalsatz der Algebra)

Jedes komplexwertige Polynom $p \in \mathbb{C}[x]$ mit $\text{Grad} > 0$ hat eine Nullstelle in $\mathbb{C}$. („ $\mathbb{C}$ ist algebraisch abgeschlossen“).

Konsequenz:

$p \in \mathbb{C}[x]$ mit $\deg(p) = n \geq 1$ zerfällt in n Linearfaktoren:

$$\begin{aligned}p(x) &= \sum_{k=0}^n a_k x^k \Rightarrow \exists x_1, \dots, x_n \in \mathbb{C} \\ p(x) &= a_n \cdot (x-x_1)(x-x_2) \cdot \dots \cdot (x-x_n)\end{aligned}$$

10.15 Beispiel:

$p(x) = 2x^2 - 3x + 5$ hat nach der abc-Formel die Nullstellen:

$$\begin{aligned}x_{1,2} &= \frac{3 \pm \sqrt{3^2 - 4 \cdot 2 \cdot 5}}{2 \cdot 2} \\ &= \frac{3 \pm \sqrt{-31}}{4} \\ &= \frac{3 \pm \sqrt{31}\sqrt{-1}}{4} \\ &= \frac{3}{4} \pm i \frac{\sqrt{31}}{4} \\ \Rightarrow p(x) &= 2 \left(x - \frac{3}{4} - i \frac{\sqrt{31}}{4} \right) \left(x - \frac{3}{4} + i \frac{\sqrt{31}}{4} \right)\end{aligned}$$

10.16 Wozu sind komplexe Zahlen noch nützlich?

- Manche mathematischen Resultate lassen sich in $\mathbb{C}$ einfacher beschreiben und besser verstehen als in $\mathbb{R}$ (z.B. das Konvergenzverhalten von Potenzreihen (später)).
- Schwingungsvorgänge in elektrischen Schaltkreisen sind in $\mathbb{C}$ wesentlich kompakter darstellbar als in $\mathbb{R}$. Es gibt also sehr konkrete Anwendungen in der Elektrotechnik und der technischen Informatik.

Kapitel 11

Polynome auf allgemeinen Körpern

11.1 Übertragung von Resultaten aus $\mathbb{R}[x]$

In Paragraph § 9 auf Seite 79 haben wir viele Resultate in $\mathbb{R}[x]$ gefunden:

- Horner-Schema
- Polynomdivision
- Abspaltung von Nullstellen
- Zahl der Nullstellen eines Polynoms n -ten Grades
- Teilbarkeit, gemeinsamer Teiler, ggT
- Euklidischer Algorithmus

Diese Resultate gelten allgemein für jeden Polynomring $K[x]$ über einem Körper K .

11.2 Beispiele

- (a) Wir betrachten das Polynom $p(x) = 4x^2 + 3x - 1$ im Körper $(\mathbb{Z}_5, +, \cdot)$.
Wegen:

$$\begin{aligned} p([0]) &= [4] \cdot [0]^2 + [3] \cdot [0] - [1] = [4] \\ p([1]) &= [4] \cdot [1]^2 + [3] \cdot [1] - [1] = [1] \\ p([2]) &= [4] \cdot [2]^2 + [3] \cdot [2] - [1] = [3] \\ p([3]) &= [4] \cdot [3]^2 + [3] \cdot [3] - [1] = [4] \\ p([4]) &= [4] \cdot [4]^2 + [3] \cdot [4] - [1] = [0] \end{aligned}$$

hat p in $\mathbb{Z}_5$ eine Nullstelle in $[4]$.

(b) Polynomdivision in $(\mathbb{Z}_3, +, \cdot)$:

+	[0]	[1]	[2]
[0]	[0]	[1]	[2]
[1]	[1]	[2]	[0]
[2]	[2]	[0]	[1]

·	[0]	[1]	[2]
[0]	[0]	[0]	[0]
[1]	[0]	[1]	[2]
[2]	[0]	[2]	[1]

$$\begin{array}{r}
 ([2]x^3 + [2]x + [1]) : (x + [2]) = [2]x^2 + [2]x + [1] \text{ Rest } [2] \\
 -([2]x^3 + x^2) \\
 \hline
 [2]x^2 + [2]x \\
 -([2]x^2 + x) \\
 \hline
 x + [1] \\
 -(x + [2]) \\
 \hline
 [2]
 \end{array}$$

Probe:

$$\begin{aligned}
 [2] + ([2]x^2 + [2]x + [1])(x + [2]) &= [2] + [2]x^3 + x^2 + [2]x^2 + x + x + [2] \\
 &= 2[x]^3 + [2]x + [1]
 \end{aligned}$$

11.3 Anwendung der Polynomdivision in $\mathbb{Z}_2[x]$ zur Datensicherung

Bei der Datenübertragung können Bits „umklappen“. Einfachste Abhilfe zur Fehlererkennung: Einführung von zusätzlichen „Prüfbits“.

Beispiel: Datenblock von 1 Byte, an den ein Prüfbit angehängt wird, das die Summe der 8 Bits modulo 2 berechnet:

$$\underbrace{11011001}_{\text{Byte}} \mid \underbrace{1}_{\text{Prüfbit}}$$

Empfängt man z.B. $11011011 \mid 1$, muss also ein Fehler aufgetreten sein.

Nachteile:

- Keine Fehlermeldung, wenn gerade Anzahl von Bits umkippt.
- Man muss $\frac{1}{8}$ mehr Daten übertragen.

Alternativ kann man Polynomdivision in $\mathbb{Z}_2[x]$ zur Datensicherung verwenden.

Vorgehensweise:

1. Interpretiere Bits eines zu übertragenden Datenblocks als Polynom $f(x) \in \mathbb{Z}_2[x]$.

Beispiel: Byte als Datenblock: 11011001

$$\text{Polynom: } f(x) = x^7 + x^6 + x^4 + x^3 + 1$$

2. Ein festes „Generatorpolynom“ $g(x) \in \mathbb{Z}_2[x]$ mit $\deg(g) = n$ dient als Divisor. Typischerweise ist $\deg(g) \ll \deg(f)$.
3. Betrachte statt $f(x)$ das Polynom $h(x) = x^n f(x)$. D.h. in der Bitfolge von $f(x)$ werden n Nullen angehängt.

Beispiel: $n = 4$

$$\begin{array}{c} \overbrace{11011001}^{h(x)} \mid \overbrace{0000}^{x^n} \\ \underbrace{11011001}_{f(x)} \end{array}$$

4. Berechne Polynomdivision $h(x) : g(x)$ in $\mathbb{Z}_2[x]$:

$$h(x) = q(x)g(x) + r(x) \quad \text{mit } \deg(r) < n.$$

5. Sende $h(x) - r(x) = q(x)g(x)$ ($= h(x) + r(x)$ in $\mathbb{Z}_2[x]$).

Wegen $\deg(r) < n$ unterscheiden sich $h(x) - r(x)$ und $h(x)$ höchstens in den letzten n Bits:

$$\begin{array}{c} \overbrace{11011001}^{h(x)-r(x)} \mid \overbrace{1011}^{r(x)} \\ \underbrace{11011001}_{f(x)} \end{array}$$

$r(x)$ dient der Datensicherung.

6. Der Empfänger erhält das Polynom $p(x)$. Tritt bei der Division von $p(x)$ durch $g(x)$ ein Rest auf, so ist $p(x) \neq h(x) - r(x)$, d.h. ein Übertragungsfehler ist aufgetreten.

Bei geschickter Wahl von $g(x)$

- ist es sehr unwahrscheinlich, dass bei einem Übertragungsfehler die Division $p(x) : g(x)$ ohne Rest aufgeht.
- kann man aus dem Divisionsrest den Fehler lokalisieren und korrigieren.

Konkrete Spezifikation im X.25 - Übertragungsprotokoll:

- Datenblock mit 4096 Byte = 32.768 Bit
 $\Rightarrow \deg(f) = 32.767$
- Generatorproblem: $g(x) = x^{16} + x^{12} + x^5 + 1$
 $\Rightarrow \deg(g) = 16$ (2 Byte zur Datensicherung)

Es werden also nur $\frac{2}{4096} \approx 0.49$ Promille zusätzliche Daten übertragen. Erkannt werden können u.A. alle 1-, 2-, 3- Bit - Fehler, sowie alle Fehler mit ungerader Fehlerzahl!

Bemerkung: Algorithmen zur Polynomdivision in $\mathbb{Z}_2[x]$ lassen sich effizient in Soft- und Hardware realisieren.

Kapitel 12

Boole'sche Algebren

12.1 Motivation

- In § 1 auf Seite 23 und § 2 auf Seite 29 haben wir gesehen, dass Operationen auf Mengen ($\cap, \cup$, Komplementbildung) und auf logischen Ausdrücken ($\wedge, \vee, \neg$) Ähnlichkeiten besitzen und ähnlichen Gesetzen folgen.
- Liegt eine gemeinsame algebraische Struktur vor?
- Gibt es weitere wichtige Beispiele hierfür?

12.2 Definition: (Boole'sche Algebra)

Eine Boole'sche Algebra besteht aus einer Menge M , zwei algebraischen Operationen $+$ und $\cdot$, sowie einer „einstelligen“ Operation $\neg$, wenn für alle $x, y, z \in M$ gilt:

(a) Kommutativgesetze:

$$\begin{aligned}x + y &= y + x \\x \cdot y &= y \cdot x\end{aligned}$$

(b) Assoziativgesetze:

$$\begin{aligned}(x + y) + z &= x + (y + z) \\(x \cdot y) \cdot z &= x \cdot (y \cdot z)\end{aligned}$$

(c) Distributivgesetze:

$$\begin{aligned}x \cdot (y + z) &= (x \cdot y) + (x \cdot z) \\x + (y \cdot z) &= (x + y) \cdot (x + z)\end{aligned}$$

(d) *Neutrale Elemente: Es gibt $0, 1 \in M$ mit*

$$0 + x = x$$

$$1 \cdot x = x$$

(e) *Komplementäre Elemente:*

$$x + \neg x = 1$$

$$x \cdot \neg x = 0$$

Wir nennen $\neg x$ das komplementäre Element zu x .

12.3 Beispiel 1: Operationen auf Mengen

Sei M eine Menge, $P(M)$ ihre Potenzmenge und $\mathbb{C}$ die Komplementbildung. Dann ist $(P(M), \cup, \cap, \mathbb{C})$ eine Boole'sche Algebra mit $\emptyset$ als Nullelement und M als Einselement.

Denn: (a), (b), (c): Kommutativ-, Assoziativ- und Distributivgesetze gelten nach Satz 1.6 auf Seite 26.

(d) Neutrale Elemente:

$$\left. \begin{array}{l} \emptyset \cup A = A \\ M \cap A = A \end{array} \right\} \quad \forall A \in P(M)$$

denn $A \subset M$.

(e) Komplementäre Elemente

Sei $A \in P(M)$ und $\mathbb{C}A := \overline{A} := M \setminus A$.

$$\Rightarrow A \cup \overline{A} = M$$

$$A \cap \overline{A} = \emptyset$$

12.4 Beispiel 2: Aussagenlogik

Die Menge der Aussagenformeln bildet zusammen mit der Disjunktion $\vee$, der Konjunktion $\wedge$ und der Negation $\neg$ eine Boole'sche Algebra. Dabei bildet die logische Konstante „falsch“ das Nullelement und „wahr“ das Einselement.

(a), (b), (c) Die Kommutativ-, Assoziativ- und Distributivgesetze gelten nach Satz 2.8 auf Seite 31.

(d) Neutrale Elemente:

$$0 \vee A = A \quad (\text{vgl. Wahrheitstafel 2.4 auf Seite 30})$$

$$1 \wedge A = A \quad (\text{vgl. Wahrheitstafel 2.4 auf Seite 30})$$

(e) Komplementäre Elemente:

$$A \vee \neg A = 1 \quad (\text{vgl. 2.8 auf Seite 31(a), Satz vom ausgeschlossenen Dritten})$$

$$A \wedge \neg A = 0 \quad (\text{vgl. 2.8 auf Seite 31(b), Satz vom Widerspruch})$$

In einer Boole'schen Algebra kann man zahlreiche Aussagen beweisen. Beispielsweise gilt:

12.5 Satz: *(Eigenschaften Boole'scher Algebren)*

In einer Boole'schen Algebra $(M, +, \cdot, \neg)$ gilt:

$$(a) \quad a = a + a, \quad a = a \cdot a \quad \forall a \in M$$

$$(b) \quad \neg(\neg a) = a \quad \forall a \in M$$

(c) Gesetze von de Morgan:

$$\neg(a + b) = (\neg a) \cdot (\neg b)$$

$$\neg(a \cdot b) = (\neg a) + (\neg b)$$

Man kann sich leicht einen Überblick über alle endlichen (!) Boole'schen Algebren verschaffen, denn es gilt:

12.6 Satz: *(Endliche Boole'sche Algebren)*

(a) Jede endliche Boole'sche Algebra ist isomorph zur Boole'schen Algebra der Teilmengen einer endlichen Menge. (vgl. 12.3 auf der vorherigen Seite).

(b) Endliche Boole'sche Algebren haben immer 2^n Elemente mit einem $n \in \mathbb{N}$

(c) Zwei endliche Algebren mit der gleichen Anzahl an Elementen sind isomorph.

Kapitel 13

Axiomatik der reellen Zahlen

13.1 Motivation

- Analysis beschäftigt sich mit Grenzwerten, Differentiation und Integration.
- Viele Phänomene in den Natur- und Ingenieurwissenschaften lassen sich mit Hilfsmitteln der Analysis beschreiben.
- In der Informatik wichtig z.B. bei Zahlendarstellung, Komplexitätsabschätzungen und physiknahen Bereichen wie Visual Computing.
- Analysis verwendet grundlegende Eigenschaften der reellen Zahlen.
- Worin unterscheidet sich $\mathbb{R}$ beispielsweise von $\mathbb{Q}$ und $\mathbb{C}$?

13.2 Definition: *(angeordnet, Positivbereich)*

Wir nennen einen Körper $(K, +, \cdot)$ angeordnet, wenn es eine Teilmenge P (den Positivbereich) gibt, so dass gilt:

- (a) $P, \{0\}$ und $-P := \{x \in K \mid -x \in P\}$ bilden eine Partition von K .
- (b) P ist abgeschlossen bzgl. $+$ und $\cdot$:

$$x, y \in P \Rightarrow x + y \in P, \quad x \cdot y \in P.$$

13.3 Definition: (Ordnungsbegriffe)

Sei $(K, +, \cdot)$ ein angeordneter Körper mit Positivbereich P . Dann lassen sich für $x, y \in K$ folgende Ordnungsbegriffe definieren:

$$\begin{aligned}x < y & :\Leftrightarrow y - x \in P \\x \leq y & :\Leftrightarrow x = y \text{ oder } x < y \\x > y & :\Leftrightarrow y < x \\x \geq y & :\Leftrightarrow y \leq x.\end{aligned}$$

Folgerung: $x \in P \Leftrightarrow x - 0 \in P \Leftrightarrow x > 0$.

Der Positivbereich enthält die positiven Zahlen.

Welche Eigenschaften besitzen angeordnete Körper?

13.4 Satz: (Eigenschaften angeordneter Körper)

Sei $(K, +, \cdot)$ ein angeordneter Körper. Dann gilt $\forall x, y, z, w \in K$:

(a) Vergleichbarkeit:

entweder $x < y$ oder $x = y$ oder $x > y$.

(b) Transitivität:

$$x < y, y < z \Rightarrow x < z.$$

(c) Verträglichkeit der Anordnung mit der Addition:

$$\text{Sei } x < y \text{ und } z \leq w : \Rightarrow x + z \leq y + w.$$

(d) Verträglichkeit der Anordnung mit der Multiplikation:

$$x < y \text{ und } z > 0 \Rightarrow xz < yz$$

$$x < y \text{ und } z < 0 \Rightarrow xz > yz.$$

(e) Invertierung:

bzgl. Addition:

$$x > 0 \Rightarrow -x < 0$$

$$x < y \Rightarrow -x > -y.$$

bzgl. Multiplikation:

$$0 < x < y \Rightarrow 0 < y^{-1} < x^{-1}.$$

(f) Positivität des Quadrats:

$$x \neq 0 \Rightarrow x^2 > 0.$$

Beweis:

Übungsaufgabe.

13.5 Konsequenzen

- (a) In einem angeordneten Körper $(K, +, \cdot)$ gilt $1 > 0$.

Beweis:

Da $(K \setminus \{0\}, \cdot)$ eine Gruppe mit neutralem Element 1 ist, ist $1 \neq 0$. Nach 13.4 (f) folgt: $0 < 1^2 = 1$.

- (b) $(\mathbb{C}, +, \cdot)$ ist kein angeordneter Körper, da $i^2 = -1 < 0$ ein Widerspruch zu 13.4 (f) ist.

- (c) Jeder angeordnete Körper enthält $\mathbb{N}$ als Teilmenge.

Denn: Aus $0 < 1$ folgt mit 13.4 (c):

$$1 < \underbrace{1+1}_{=:2} < \underbrace{1+1+1}_{=:3} < \dots$$

Endliche Körper wie z.B. $\mathbb{Z}_p$, (p Primzahl) können daher nicht angeordnet sein.

Offensichtlich sind $(\mathbb{Q}, +, \cdot)$ und $(\mathbb{R}, +, \cdot)$ angeordnete Körper.

Worin unterscheiden sich $\mathbb{R}$ und $\mathbb{Q}$?

13.6 Definition: (*Maximum und Minimum*)

Sei $(K, +, \cdot)$ ein angeordneter Körper und $M \subset K$.

Ein Element $x \in M$ heißt Maximum von M ($x = \max M$), wenn für alle $y \in M$ gilt: $y \leq x$.

$x \in M$ heißt Minimum von M ($x = \min M$), falls $y \geq x$ für alle $y \in M$

Bemerkung: Hat M ein Maximum oder Minimum, so ist dieses eindeutig bestimmt.

Denn: Seien x, y zwei Maxima. Dann gilt:

$$\left. \begin{array}{l} y \geq x \quad (\text{da } y \text{ Maximum}) \\ x \geq y \quad (\text{da } x \text{ Maximum}) \end{array} \right\} \stackrel{13.7(a)}{\Rightarrow} x = y.$$

□

13.7 Beispiel:

Sei $K = \mathbb{Q}$.

- (a) $M_1 := \{x \in \mathbb{Q} \mid x \leq 2\} \Rightarrow \max M_1 = 2$, da $2 \in M_1$ und $x \leq 2 \forall x \in M_1$.

- (b) $M_2 := \{x \in \mathbb{Q} \mid x < 2\}$ hat kein Maximum, denn zu jedem $x \in M_2$ gibt es ein $y \in M_2$ mit $y > x$. Wähle z.B. $y := \frac{x+2}{2}$. Dann ist $x < y < 2$.
- (c) $M_3 := \{x \in \mathbb{Q} \mid x^2 \leq 2\}$ hat kein Maximum, denn: Ist x Maximum, so ist $x^2 = 2$ oder $x^2 < 2$. Für $x^2 = 2$ ist $x \notin \mathbb{Q}$. Für $x^2 < 2$ findet man ähnlich wie in (b) eine Zahl $y \in \mathbb{Q}$ mit $x < y$ und $y^2 < 2$.

Gibt es ein „Maximum, das nicht in M liegt“?

13.8 Definition: (*Infimum, Supremum*)

Sei $(K, +, \cdot)$ ein angeordneter Körper und $M \subset K$.

M heißt nach oben beschränkt, wenn es ein $x \in K$ (nicht notwendigerweise in M) gibt mit $y \leq x \forall y \in M$. Dann heißt x obere Schranke von M .

$s \in K$ heißt kleinste obere Schranke (Supremum) von M ($s = \sup M$), wenn für jede obere Schranke x gilt: $x \geq s$.

Analog definiert man nach unten beschränkt, untere Schranke und größte untere Schranke (Infimum) ($r = \inf M$).

13.9 Beispiel:

Wir betrachten $K = \mathbb{Q}$ und Beispiele aus 13.7 auf der vorherigen Seite.

- (a) M_1 hat viele obere Schranken, z.B. 1000; 7; 2.
Kleinste obere Schranke: $\sup M_1 = 2$.
 M_1 ist nicht nach unten beschränkt.
- (b) M_2 hat die gleichen oberen Schranken wie M_1 .
Obwohl kein Maximum existiert, gibt es ein Supremum: $\sup M_2 = 2$.
- (c) M_3 hat viele obere Schranken, z.B. 1000; 7; 1,5.
Es gibt jedoch keine kleinste obere Schranke in $\mathbb{Q}$, da $\sqrt{2}$ keine rationale Zahl ist.

Diese Beispiele illustrieren:

- (a) Wenn eine Menge M ein Maximum hat, so ist dieses auch das Supremum.
- (b) Es gibt Mengen, die ein Supremum besitzen, jedoch kein Maximum.
- (c) In $\mathbb{Q}$ gibt es Mengen, die obere Schranken besitzen, jedoch keine kleinste obere Schranke.

13.10 Definition: (vollständig)

Ein angeordneter Körper heißt vollständig, wenn in ihm jede nach oben beschränkte Menge ein Supremum besitzt.

Man kann zeigen:

13.11 Satz: (Axiomatische Charakterisierung von $\mathbb{R}$)

Es gibt genau einen vollständigen, angeordneten Körper.

Er heißt Körper der reellen Zahlen $\mathbb{R}$.

Mit der Unbeschränktheit der natürlichen Zahlen zeigt man:

13.12 Satz: (Eigenschaften der reellen Zahlen)

Für $x, y \in \mathbb{R}$ gilt:

- (a) Zu $x, y > 0$ existiert $n \in \mathbb{N}$ mit $nx > y$ (Archimedisches Axiom).
- (b) Zu $x > 0$ existiert $n \in \mathbb{N}$ mit $\frac{1}{n} < x$.
- (c) Zu $x \in \mathbb{R}$ existiert $m := \max\{z \in \mathbb{Z} \mid z \leq x\}$.

13.13 Definition: (Betragsfunktion)

Für $x \in \mathbb{R}$ heißt $|x| := \begin{cases} x & \text{falls } x \geq 0 \\ -x & \text{falls } x < 0 \end{cases}$

Betrag von x .

13.14 Satz: (Eigenschaften des Betrags)

Für $x, y \in \mathbb{R}$ gilt:

- (a) $|x| \geq 0$
 $|x| = 0 \Leftrightarrow x = 0$
- (b) $|x \cdot y| = |x| \cdot |y|$
- (c) $|x + y| \leq |x| + |y|$
- (d) $\left| \frac{x}{y} \right| = \frac{|x|}{|y|}$

$$(e) \quad ||x| - |y|| \leq |x - y|$$

Teil III

Eindimensionale Analysis

Kapitel 14

Folgen

14.1 Motivation

- In der Mathematik werden Folgen benötigt, um stetige Funktionen zu untersuchen und um die Differential- und Integralrechnung einzuführen.
- Klassische Funktionen wie $\sin$, $\cos$, $\exp$, $\log$ werden als Grenzwerte von Zahlenfolgen definiert. Programmiert man eine Mathematikbibliothek, ist es nötig, Darstellungen zu verwenden, bei denen die Folgen schnell berechenbar sind und schnell konvergieren.

14.2 Definition: (Folge)

Eine Abbildung $f : \mathbb{N} \mapsto \mathbb{R}, n \mapsto f(n) =: a_n$ heißt (reellwertige) Folge. Wir nennen a_n das n -te Glied der Folge und kürzen die Folge mit $(a_n)_{n \in \mathbb{N}}$, (a_n) oder einfach nur a_n ab.

Bemerkung: Viele der nachfolgenden Resultate (aber nicht alle) lassen sich verallgemeinern auf Folgen mit anderen Wertebereichen, z.B. $\mathbb{C}, \mathbb{R}^n$.

14.3 Beispiel:

- (a) konstante Folge: $a_n = a \quad \forall n \in \mathbb{N}$.
- (b) $\left(\frac{1}{n}\right)_{n \in \mathbb{N}}$ ist die Folge $1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots$
- (c) $\left(\frac{n}{n+1}\right)_{n \in \mathbb{N}}$ ergibt $\frac{1}{2}, \frac{2}{3}, \frac{3}{4}, \dots$

(d) Die Fibonacci-Folge ist rekursiv definiert durch

$$\begin{aligned}a_1 &:= 1, & a_2 &:= 1 \\a_n &:= a_{n-1} + a_{n-2} & \text{für } n \geq 3\end{aligned}$$

Sie ergibt: 1, 1, 2, 3, 5, 8, 13, 21, ...

(e) $(b^n)_{n \in \mathbb{N}}$ ist die Folge: $b, b^2, b^3, \dots$

14.4 Definition: ((streng) monoton fallend / wachsend; beschränkt)

Eine reellwertige Folge (a_n) heißt monoton wachsend, wenn $a_{n+1} \geq a_n \quad \forall n \in \mathbb{N}$ ist. Gilt sogar $a_{n+1} > a_n \quad \forall n \in \mathbb{N}$, so ist (a_n) streng monoton wachsend. Analog definiert man (streng) monoton fallend.

(a_n) heißt nach oben/unten beschränkt, wenn $\{(a_n) \mid n \in \mathbb{N}\}$ nach oben/unten beschränkt ist.

$\sup_{n \in \mathbb{N}} a_n$ und $\inf_{n \in \mathbb{N}} a_n$ werden als Supremum/Infimum von $\{(a_n) \mid n \in \mathbb{N}\}$ definiert.

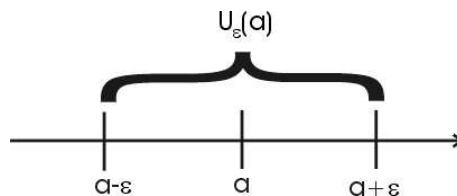
Wie kann man das Konvergenzverhalten einer Folge beschreiben?

14.5 Definition: (ε - Umgebung)

Sei $a \in \mathbb{R}$ und $\varepsilon > 0$ eine reelle Zahl. Dann nennen wir

$$U_\varepsilon(a) := \{x \in \mathbb{R} \mid |x - a| < \varepsilon\} = \{x \in \mathbb{R} \mid a - \varepsilon < x < a + \varepsilon\}$$

die ε - Umgebung von a .



14.6 Definition: (konvergente Folge)

Eine Folge (a_n) heißt konvergent gegen a , wenn in jeder ε - Umgebung von a „fast alle“ Folgenglieder liegen:

$$\forall \varepsilon > 0 \quad \exists n_0(\varepsilon) \in \mathbb{N} \quad \text{mit} \quad |x - a| < \varepsilon \quad \forall n \geq n_0(\varepsilon).$$

Man schreibt „ $\lim_{n \rightarrow \infty} a_n = a$ “ oder „ $a_n \rightarrow a$ für $n \rightarrow \infty$ “ und nennt a den Grenzwert (Limes) von (a_n) .

Eine reellwertige Folge heißt divergent, wenn sie gegen keine reelle Zahl konvergiert.

14.7 Beispiel:

- (a) Die konstante Folge ist konvergent:

$$a_n = a \quad \forall n \in \mathbb{N} \Rightarrow \lim_{n \rightarrow \infty} a_n = a.$$

- (b) $\lim_{n \rightarrow \infty} \frac{1}{n} = 0$. Denn :

Sei $\varepsilon > 0$. Dann existiert nach 13.12 (b) ein $n_0(\varepsilon) \in \mathbb{N}$ mit $\frac{1}{n_0} < \varepsilon$.

Wegen $0 < \frac{1}{n} \leq \frac{1}{n_0} \quad \forall n \geq n_0$ ist also $\frac{1}{n} \in U_\varepsilon(0) \quad \forall n \geq n_0$.

□

Folgen, die gegen 0 konvergieren, heißen Nullfolgen.

- (c) Die Folge $((-1)^n)_{n \in \mathbb{N}}$ ist divergent.

14.8 Definition: (uneigentliche Konvergenz, bestimmte Divergenz)

Sei (a_n) eine Folge. Dann strebt a_n gegen ∞ , ($\lim_{n \rightarrow \infty} a_n = \infty$), falls für jedes $r > 0$ ein $n_0(r) \in \mathbb{N}$ existiert mit $a_n > r \quad \forall n \geq n_0$. In diesem Fall spricht man auch von uneigentlicher Konvergenz oder bestimmter Divergenz.

Analog definiert man $\lim_{n \rightarrow \infty} a_n = -\infty$, falls für jedes $r < 0$ ein $n_0(r) \in \mathbb{N}$ existiert mit $a_n < r \quad \forall n \geq n_0$.

Beispiel: $\lim_{n \rightarrow \infty} n^2 = \infty$

Können unbeschränkte Folgen konvergieren?

14.9 Satz: (Beschränktheit konvergenter Folgen)

Eine konvergente Folge ist beschränkt, d.h. sie ist sowohl nach oben als auch nach unten beschränkt.

Beweis:

Sei (a_n) Folge mit $\lim_{n \rightarrow \infty} a_n = a$. Dann existiert n_0 mit $a_n \in U_1(a) \quad \forall n \geq n_0$. Damit ist $\{a_n \mid n \in \mathbb{N}\}$ in der beschränkten Menge $\{a_1, a_2, \dots, a_{n_0-1}\} \cup U_1(a)$ enthalten.

□

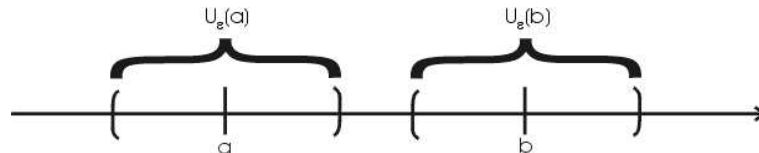
Kann es mehrere Grenzwerte geben?

14.10 Satz: (Eindeutigkeit des Grenzwerts)

Der Grenzwert einer konvergenten Folge ist eindeutig.

Beweis:

Annahme: (a_n) habe zwei Grenzwerte a, b mit $a \neq b$.



Wähle $\varepsilon > 0$ so, dass $\varepsilon < \frac{|a-b|}{2}$. Dann ist

$$U_\varepsilon(a) \cap U_\varepsilon(b) = \emptyset \quad (*)$$

Da a, b Grenzwerte sind, existiert n_0, n_1 mit

$$\begin{aligned} a_n &\in U_\varepsilon(a) & \forall n \geq n_0 \text{ und} \\ a_n &\in U_\varepsilon(b) & \forall n \geq n_1 \end{aligned}$$

D.h. $a_n \in U_\varepsilon(a) \cap U_\varepsilon(b) \quad \forall n \geq \max(n_0, n_1)$.

Das ist ein Widerspruch zu $(*)$. $\Rightarrow a = b$.

□

Gibt es Kriterien, mit deren Hilfe man Konvergenz zeigen kann?

14.11 Satz: (Konvergenzkriterien)

(a) Vergleichskriterium: Seien $(a_n), (b_n), (c_n)$ reelle Folgen mit $a_n \leq b_n \leq c_n \quad \forall n \in \mathbb{N}$ und $\lim_{n \rightarrow \infty} a_n = b = \lim_{n \rightarrow \infty} c_n$. Dann konvergiert (b_n) und es gilt $\lim_{n \rightarrow \infty} b_n = b$.

(b) Cauchy-Kriterium: Eine reelle Zahlenfolge ist genau dann konvergent, wenn sie eine Cauchy-Folge ist, d.h.

$$\forall \varepsilon > 0 \exists n_0(\varepsilon) \in \mathbb{N} : |a_n - a_m| < \varepsilon \quad \forall n, m \geq n_0(\varepsilon).$$

(d.h. die Folgenglieder weichen beliebig wenig voneinander ab, wenn der Index hinreichend groß ist.)

(c) Eine monoton wachsende, nach oben beschränkte Folge konvergiert. Ebenso konvergiert eine monoton fallende, nach unten beschränkte Folge.

Beweis:

von (a)

Sei $\varepsilon > 0$. Dann existieren $n_0, n_1 \in \mathbb{N}$:

$$\begin{aligned} a_n &\in U_\varepsilon(b) & \forall n \geq n_0 \\ c_n &\in U_\varepsilon(b) & \forall n \geq n_1 \end{aligned}$$

$$\Rightarrow b - \varepsilon < a_n \leq b_n \leq c_n < b + \varepsilon \quad \forall n \geq \max(n_0, n_1) =: n_2$$

Also ist $b_n \in U_\varepsilon(b) \quad \forall n \geq n_0$, d.h. $\lim_{n \rightarrow \infty} b_n = b$.

□

14.12 Beispiel:

$\left(\frac{1}{n^2}\right)$ ist monoton fallend und von unten durch 0 beschränkt. Also ist $\left(\frac{1}{n^2}\right)$ nach 14.11 (c) konvergent.

14.13 Satz: (Rechenregeln für Grenzwerte)

Seien $(a_n), (b_n)$ reelle Folgen mit $\lim_{n \rightarrow \infty} a_n = a, \lim_{n \rightarrow \infty} b_n = b$. Dann gilt:

- (a) Falls $a_n \leq b_n \quad \forall n \in \mathbb{N}$, so ist $a \leq b$.
- (b) Falls $a_n < b_n \quad \forall n \in \mathbb{N}$, so ist $a \leq b$ ($a < b$ ist i.A. falsch!).
- (c) $\lim_{n \rightarrow \infty} (a_n \pm b_n) = a \pm b$
- (d) $\lim_{n \rightarrow \infty} (a_n \cdot b_n) = a \cdot b$ (d.h. insbesondere $\lim_{n \rightarrow \infty} (c \cdot a_n) = c \cdot a$).
- (e) Ist $b \neq 0$, so existiert n_0 mit $b_n \neq 0 \quad \forall n \geq n_0$. Dann sind auch $\left(\frac{1}{b_n}\right)_{n \geq n_0}, \left(\frac{a_n}{b_n}\right)_{n \geq n_0}$ konvergent und haben den Limes $\frac{1}{b}$ bzw. $\frac{a}{b}$.
- (f) $(|a_n|)$ konvergiert gegen $|a|$.
- (g) Sei $m \in \mathbb{N}$ und $a_n \geq 0 \quad \forall n \in \mathbb{N} \Rightarrow \lim_{n \rightarrow \infty} \sqrt[m]{a_n} = \sqrt[m]{\lim_{n \rightarrow \infty} a_n}$
Dabei bezeichnet $\sqrt[m]{a_n}$ die eindeutig bestimmte Zahl $w \geq 0$ mit $w^m = a_n$.

14.14 Beispiel:

$$(a) \quad \lim_{n \rightarrow \infty} \frac{17}{n} = 17 \lim_{n \rightarrow \infty} \frac{1}{n} = 17 \cdot 0 = 0.$$

(b)

$$\begin{aligned}\lim_{n \rightarrow \infty} \frac{n}{n+1} &= \lim_{n \rightarrow \infty} \left(\frac{n}{n} \cdot \frac{1}{1 + \frac{1}{n}} \right) \\ &= \lim_{n \rightarrow \infty} 1 \cdot \frac{\lim_{n \rightarrow \infty} 1}{\lim_{n \rightarrow \infty} 1 + \lim_{n \rightarrow \infty} \frac{1}{n}} \\ &= 1 \cdot \frac{1}{1+0} = 1\end{aligned}$$

$$(c) \quad \lim_{n \rightarrow \infty} \frac{5n^4 - 2n^2 + 1}{7n^4 + 11n^3 + 1} = \frac{5 - \frac{2}{n^2} + \frac{1}{n^4}}{7 + \frac{11}{n} + \frac{1}{n^4}} = \frac{5 - 0 + 0}{7 + 0 + 0} = \frac{5}{7}.$$

14.15 Der Grenzwert $\lim_{n \rightarrow \infty} \left(1 + \frac{p}{n}\right)^n$

Anwendungsbeispiel:

Ein Geldbetrag a_0 wird mit einem Jahreszinsfuß p jährlich, halbjährlich, vierteljährlich, monatlich oder täglich verzinst. Nach einem Jahr ergibt sich also:

$$\begin{aligned}a_1 &= a_0 (1 + p) && \text{bei jährlicher Verzinsung} \\ a_2 &= a_0 \left(1 + \frac{p}{2}\right)^2 && \text{bei halbjährlicher Verzinsung} \\ a_4 &= a_0 \left(1 + \frac{p}{4}\right)^4 && \text{bei vierteljährlicher Verzinsung} \\ a_{12} &= a_0 \left(1 + \frac{p}{12}\right)^{12} && \text{bei monatlicher Verzinsung} \\ a_{365} &= a_0 \left(1 + \frac{p}{365}\right)^{365} && \text{bei täglicher Verzinsung}\end{aligned}$$

Z.B. ergibt mit $a_0 = 100$ Euro und $p = 0,03$:

$$\begin{aligned}a_1 &= 103 \text{ Euro} \\ a_2 &= 103.02 \text{ Euro} \\ a_4 &= 103.03 \text{ Euro} \\ a_{12} &= 103.04 \text{ Euro} \\ a_{365} &= 103.05 \text{ Euro}\end{aligned}$$

Konvergiert die Folge $\left(1 + \frac{p}{n}\right)^n$?

Man kann zeigen :

- (a_n) ist streng monoton wachsend.
- (a_n) ist nach oben beschränkt.

Nach Satz 14.11 (c) konvergiert damit (a_n) .
Für $p = 1$ ist der Grenzwert die Eulersche Zahl **e**.

$$\lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n = e \approx 2,7182\dots$$

und für ein allgemeines p erhält man

$$\boxed{\lim_{n \rightarrow \infty} \left(1 + \frac{p}{n}\right)^n = e^p}$$

Kapitel 15

Landau - Symbole

15.1 Motivation

- Suchen kompakte Notation zur Aufwandsabschätzung von Algorithmen.
- Algorithmus soll von Problemgröße n abhängen (z.B.: n = Zahl der zu sortierenden Objekte in einer Liste)
- Laufzeit ist Abbildung $f : \mathbb{N} \mapsto \mathbb{R}^+$ mit $f(n) = a_n$, d.h. eine Folge.
- Meist ist (a_n) monoton wachsend und nicht beschränkt, d.h. $\lim_{n \rightarrow \infty} a_n = \infty$.
- Interessant ist, wie schnell a_n gegen ∞ strebt. Hierzu vergleicht man (a_n) mit Prototypen von anderen Folgen (b_n) .

15.2 Definition: (Groß-O, Klein-o)

Ein Folge $A = (a_n)$ ist „Groß-O“ von $B = (b_n)$, wenn der Quotient $\left(\frac{a_n}{b_n}\right)$ beschränkt ist. Die Folge A ist „Klein-o“ von B , wenn $\left(\frac{a_n}{b_n}\right)$ eine Nullfolge ist. Wir schreiben $A = O(B)$ bzw. $A = o(B)$. O und o nennt man auch Landau-Symbole.

15.3 Beispiel:

(a) $2n^2 + 3n + 4 = O(n^2)$, denn

$$\frac{2n^2 + 3n + 4}{n^2} = \frac{n^2(2 + \frac{3}{n} + \frac{4}{n^2})}{n^2} = 2 + \frac{3}{n} + \frac{4}{n^2}$$

Dies ist eine konvergente Folge und damit beschränkt.

(b) $2n^2 + 3n + 4 = o(n^3)$, denn

$$\frac{2n^2 + 3n + 4}{n^3} = \frac{\frac{2}{n} + \frac{3}{n^2} + \frac{4}{n^3}}{1} \rightarrow 0 \quad \text{für } n \rightarrow \infty.$$

15.4 Mehrdeutigkeit

Die Angabe $A = O(B)$ ist nicht eindeutig. Für $A = (n + 2)$ sind z.B.:

$$\begin{aligned} A &= O(n^2), \\ A &= O(n), \\ A &= O\left(\frac{n}{2}\right) \end{aligned}$$

wahre Aussagen. Es gilt:

- Konstante Faktoren ändern nicht die Ordnung.

$$O(B) = O(c \cdot B) \quad \forall c \in \mathbb{R}.$$

- Terme niedriger Ordnung sind unwichtig:

z.B. haben $(n^2 + 3n + 4)$ und $(n^2 + 17n)$ die selbe Ordnung $O(n^2)$.

15.5 Satz: (Rechenregeln für Landau-Symbole)

Für reelle Zahlenfolgen A , B und C gelten folgen Regeln:

- (a) $A = O(A)$
- (b) $\left. \begin{aligned} c \cdot O(A) &= O(A) \\ c \cdot o(A) &= o(A) \end{aligned} \right\} \forall c \in \mathbb{R}$
- (c) $\begin{aligned} O(A) + O(A) &= O(A) \\ o(A) + o(A) &= o(A) \end{aligned}$
- (d) $\begin{aligned} O(O(A)) &= O(A) \\ o(o(A)) &= o(A) \end{aligned}$
- (e) $\begin{aligned} O(A) \cdot O(B) &= O(A \cdot B) \\ o(A) \cdot o(B) &= o(A \cdot B) \end{aligned}$
- (f) $\begin{aligned} A \cdot O(B) &= O(A \cdot B) \\ A \cdot o(B) &= o(A \cdot B) \end{aligned}$
- (g) *Transitivität:*
- $$\begin{aligned} A = O(B), \quad B = O(C) &\Rightarrow A = O(C) \\ A = o(B), \quad B = o(C) &\Rightarrow A = o(C) \end{aligned}$$

15.6 Vergleichbarkeit von Folgen

Nicht alle Folgen sind vergleichbar. Betrachte z.B. $A = (a_n)$, $B = (b_n)$ mit

$$a_n = \begin{cases} n^2 & (n \text{ gerade}) \\ n & (n \text{ ungerade}) \end{cases}$$
$$b_n = \begin{cases} n & (n \text{ gerade}) \\ n^2 & (n \text{ ungerade}) \end{cases}$$

Für gerades n sieht man, dass $A \neq O(B)$.

Für ungerades n sieht man, dass $B \neq O(A)$.

15.7 Definition: $(O(A) = O(B), O(A) < O(B))$

Wir sagen

$$O(A) = O(B) \quad :\Leftrightarrow \quad A = O(B) \text{ und } B = O(A)$$

$$O(A) < O(B) \quad :\Leftrightarrow \quad A = O(B) \text{ und } B \neq O(A)$$

15.8 Häufig verwendete Prototypen von Vergleichsfunktionen

O	Laufzeitverhalten
$O(1)$	konstant
$O(\log_a n)$, $a > 1$	logarithmisch
$O(n)$	linear
$O(n \log_a n)$, $a > 1$	$n \log n$
$O(n^2)$	quadratisch
$O(n^3)$	kubisch
$O(n^k)$	polynomial
$O(a^n)$, $a > 1$	exponentiell

Häufig ist $a = 2$. Die Wahl von a ist jedoch unwichtig, da sich Logarithmen zu unterschiedlichen Basen nur durch eine multiplikative Konstante unterscheiden.

$$\log_a n = \frac{\text{ld } n}{\text{ld } a} \quad \text{mit } \text{ld} := \log_2.$$

15.9 Vergleich von Ordnungen

Es gilt:

$$\begin{array}{ccccccc} O(1) & < O(\log n) & < O(n) & < O(n \log n) & < O(n^2) & < O(n^3) \\ & < \dots & < O(n^k) & < \dots & < O(2^n) & < O(3^n) & < \dots \end{array}$$

n	$\text{ld } n$	$n \text{ ld } n$	n^2	n^3	2^n
10	3,32	33,22	100	1000	1024
100	6,64	664,4	10000	10^6	$1,27 \cdot 10^{30}$
1000	9,97	9966	10^6	10^9	$\approx 10^{301}$
10000	13,29	132877	10^8	10^{12}	$\approx 10^{3010}$

Für große n sollte man daher nicht auf Fortschritte im Rechnerbau hoffen, sondern nach Algorithmen suchen, die eine bessere Ordnung besitzen.

Kapitel 16

Reihen

16.1 Motivation

- Reihen sind spezielle Folgen, die z.B. in Form von Potenzreihen wichtige Anwendungen in der Approximation klassischer Funktionen wie $\sin$, $\cos$, $\exp$, $\log$ haben.
- In der Informatik haben sie zudem Bedeutung in der Zahlendarstellung.

16.2 Definition: (Reihe, n -te Partialsumme)

Sei (a_n) eine Folge. Bildet man hieraus eine neue Folge (s_n) mit

$$s_n := \sum_{k=1}^n a_k, \quad n \in \mathbb{N},$$

so nennt man (s_n) eine Reihe und schreibt hierfür $\sum_{k=1}^{\infty} a_k$. Das n -te Glied s_n

heißt n -te Partialsumme der Reihe. Ist die Reihe $\sum_{k=1}^{\infty} a_k$ konvergent, so wird ihr

Grenzwert $s := \lim_{n \rightarrow \infty} \sum_{k=1}^n a_k$ ebenfalls mit $\sum_{k=1}^{\infty} a_k$ bezeichnet.

Folgen und Reihen unterscheiden sich also lediglich darin, dass man bei Reihen versucht, Konvergenzaussagen nicht in Abhängigkeit von s_n , sondern von den Folgengliedern a_k zu gewinnen.

16.3 Satz: (Konvergenzkriterien für Reihen)

(a) Cauchy-Kriterium: (vgl. 14.11 (b))

$\sum_{k=1}^{\infty} a_k$ konvergiert genau dann, wenn es zu jedem $\epsilon > 0$ ein $n_0(\epsilon) \in \mathbb{N}$ gibt mit

$$\left| \sum_{k=n}^m a_k \right| < \epsilon \quad \forall m, n \geq n_0(\epsilon) \text{ mit } m \geq n.$$

(b) Notwendige Konvergenzbedingung:

$\sum_{k=1}^{\infty} a_k$ konvergent $\Rightarrow (a_k)$ ist Nullfolge.

(Die Umkehrung gilt **nicht!**)

(c) Linearität

Falls $\sum_{k=1}^{\infty} a_k, \sum_{k=1}^{\infty} b_k$ konvergieren, so konvergieren auch $\sum_{k=1}^{\infty} (a_k \pm b_k)$ und

$\sum_{k=1}^{\infty} (c \cdot a_k)$. Es gilt:

$$\begin{aligned} \sum_{k=1}^{\infty} (a_k \pm b_k) &= \sum_{k=1}^{\infty} a_k \pm \sum_{k=1}^{\infty} b_k \\ \sum_{k=1}^{\infty} (c \cdot a_k) &= c \cdot \sum_{k=1}^{\infty} a_k \end{aligned}$$

(d) Leibniz-Kriterium:

Sei (a_k) eine monoton fallende Nullfolge nichtnegativer Zahlen a_k . Dann konvergieren die alternierenden Reihen $\sum_{k=1}^{\infty} (-1)^k a_k$ und $\sum_{k=1}^{\infty} (-1)^{k+1} a_k$.

Der Grenzwert liegt zwischen zwei aufeinander folgenden Partialsummen.

Beweis:

(a) folgt aus dem Cauchy-Kriterium 14.11 (b).

(b) folgt aus (a) mit $m = n$.

(c) folgt aus Satz 14.13 auf Seite 115.

(d) Betrachte (s_n) mit $s_n := \sum_{k=1}^n (-1)^k a_k$. Dann gilt:

$$s_{2n+2} - s_{2n} = a_{2n+2} - a_{2n+1} \leq 0 \quad \text{d.h. } (s_{2n})_{n \in \mathbb{N}} \text{ monoton fallend. (*)}$$

$$s_{2n+1} - s_{2n-1} = -a_{2n+1} + a_{2n} \geq 0 \quad \text{d.h. } (s_{2n-1})_{n \in \mathbb{N}} \text{ monoton wachsend. (**)}$$

Wegen

$$s_{2n} = s_{2n-1} + a_{2n} \geq s_{2n-1} \stackrel{\text{wachsend}}{\geq} s_1 \quad \forall n \in \mathbb{N}$$

ist (s_{2n}) nach unten beschränkt. Damit existiert

$$s_n := \lim_{n \rightarrow \infty} s_{2n}.$$

Analog zeigt man, dass (s_{2n-1}) nach oben beschränkt ist und somit

$$s' := \lim_{n \rightarrow \infty} s_{2n-1}$$

existiert.

Wegen

$$s - s' = \lim_{n \rightarrow \infty} a_{2n} = 0$$

haben (s_{2n}) und (s_{2n-1}) den selben Grenzwert s . Also konvergiert (s_n) gegen s .

Ferner gilt:

$$s_{2n-1} \stackrel{(**)}{\leq} s \stackrel{(*)}{\leq} s_{2n} \quad \forall n.$$

Der Beweis für $s_n := \sum_{k=1}^{\infty} (-1)^{k+1} a_k$ verläuft analog.

□

Bemerkung: Das Cauchy-Kriterium impliziert, dass das Konvergenzverhalten einer Reihe sich nicht ändert, wenn man endlich viele Folgenglieder abändert. In diesem Fall ändert sich jedoch i. A. der Grenzwert.

16.4 Beispiel:

(a) Geometrische Reihe:

$$\sum_{k=0}^{\infty} q^k = 1 + q + q^2 + q^3 + \dots$$

(Hier ist es sinnvoll mit $k = 0$ zu beginnen.)

Wie sehen die Partialsummen aus?

$$\begin{array}{rcl} s_n & = & 1 + q + q^2 + q^3 + \dots + q^n \\ q \cdot s_n & = & q + q^2 + q^3 + \dots + q^n + q^{n+1} \\ \hline (1-q) \cdot s_n & = & 1 - q^{n+1} \end{array}$$

$$\Rightarrow \text{Für } q \neq 1: \quad s_n = \frac{1-q^{n+1}}{1-q}$$

Für $|q| < 1$ ist daher

$$\sum_{k=0}^{\infty} q^k = \lim_{n \rightarrow \infty} s_n = \frac{1}{1-q}$$

d.h. die geometrische Reihe konvergiert.

Für $|q| \geq 1$ ist $(a_k) = (q^k)$ keine Nullfolge. Nach 16.3 (b) kann $\sum_{k=0}^{\infty} q^k$ in diesem Fall nicht konvergieren.

(b) Harmonische Reihe:

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

ist divergent, denn es gilt die Abschätzung:

$$\sum_{k=n}^m \frac{1}{k} \geq \sum_{k=n}^m \frac{1}{m} = \frac{m-n+1}{m} \rightarrow 1 \text{ für } m \rightarrow \infty$$

Somit ist für $0 < \epsilon < 1$ das Cauchy-Kriterium verletzt.

(c) Alternierende harmonische Reihe:

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

konvergiert nach dem Leibniz-Kriterium, da $\left(\frac{1}{n}\right)_{n \in \mathbb{N}}$ eine monoton fallende Nullfolge ist.

Wegen $s_1 = 1$, $s_2 = \frac{1}{2}$ gilt die Einschließung

$$\frac{1}{2} \leq \sum_{k=1}^{\infty} (-1)^{k+1} \frac{1}{2} \leq 1.$$

(Man kann zeigen, dass der Grenzwert $\ln 2 \approx 0,69$ beträgt.)

Bei endlichen Summen darf man die Reihenfolge der Summation beliebig vertauschen. Bei unendlichen Summen (Reihen) ist diese nur in besonderen Fällen erlaubt. Hierzu benötigen wir den Begriff der absolut konvergenten Reihe.

16.5 Definition: (absolut konvergente Reihe)

Eine Reihe $\sum_{k=1}^{\infty} a_k$ heißt absolut konvergent, wenn die Reihe $\sum_{k=1}^{\infty} |a_k|$ konvergiert.

16.6 Beispiel:

- (a) Jede absolut konvergente Reihe ist konvergent.

Beweis:

Sei $\sum_{k=1}^{\infty} a_k$ absolut konvergent. Nach dem Cauchy-Kriterium für die Reihe $\sum_{k=1}^{\infty} |a_k|$ existiert zu jedem $\epsilon > 0$ ein $n_0(\epsilon) \in \mathbb{N}$ mit

$$\sum_{k=n}^m |a_k| < \epsilon \quad \forall n, m \geq n_0 \text{ mit } m \geq n$$

Wegen $\left| \sum_{k=n}^m a_k \right| \leq \sum_{k=n}^m |a_k|$ (vgl. 13.14 (c)) gilt:

$$\left| \sum_{k=n}^m a_k \right| < \epsilon \quad \forall n, m \geq n_0 \text{ mit } m \geq n$$

d.h. $\sum_{k=1}^{\infty} a_k$ konvergiert nach dem Cauchy-Kriterium.

- (b) Es gibt konvergente Reihen, die nicht absolut konvergieren.

Beispiel: Die alternierende harmonische Reihe (16.4 (c)) konvergiert. Sie konvergiert jedoch nicht absolut, da die harmonische Reihe (16.4 (b)) divergiert.

16.7 Satz: (Kriterien für absolute Konvergenz)

- (a) Die Reihe $\sum_{k=1}^{\infty} a_k$ ist absolut konvergent.

$\Leftrightarrow$ Die Folge $\left(\sum_{k=1}^n |a_k| \right)_{n \in \mathbb{N}}$ ist beschränkt.

- (b) Majorantenkriterium:

$|a_k| \leq b_k$ für alle $k \in \mathbb{N}$ und $\sum_{k=1}^{\infty} b_k$ konvergiert.

$\Rightarrow \sum_{k=1}^{\infty} a_k$ ist absolut konvergent.

(Die Reihe $\sum_{k=1}^{\infty} b_k$ heißt Majorante für $\sum_{k=1}^{\infty} a_k$.)

(c) Quotientenkriterium:

Es sei $a_k \neq 0$ für alle $k \geq k_0$, $k_0 \in \mathbb{N}$. Gilt für alle $k \geq k_0$ die Ungleichung $\left| \frac{a_{k+1}}{a_k} \right| \leq q$ mit einem $q < 1$, so ist $\sum_{k=1}^{\infty} a_k$ absolut konvergent.

(d) Wurzelkriterium:

Gilt für alle $k \geq k_0$ die Ungleichung $\sqrt[k]{|a_k|} \leq q$ mit einem $q < 1$, so ist $\sum_{k=1}^{\infty} a_k$ absolut konvergent.

Beweis:

(a) $\left(\sum_{k=1}^n |a_k| \right)_{n \in \mathbb{N}}$ ist monoton wachsende Folge. Dann ist sie nach 14.11(c) konvergent, falls sie beschränkt ist („ $\Rightarrow$ “). Die andere Richtung („ $\Leftarrow$ “) folgt aus Satz 14.9.

□

(b) Die Folge $\left(\sum_{k=1}^n b_k \right)_{n \in \mathbb{N}}$ der Partialsummen von $\sum_{k=1}^{\infty} b_k$ ist konvergent und damit beschränkt.

Wegen $\sum_{k=1}^n |a_k| \leq \sum_{k=1}^n b_k$ ist die Folge der Partialsummen von $\sum_{k=1}^{\infty} |a_k|$ durch dieselbe Schranke nach oben beschränkt. Nach (a) folgt absolute Konvergenz von $\sum_{k=1}^{\infty} a_k$.

□

(c) und (d) In beiden Fällen wird der Beweis geführt, indem die Reihe gemäß Majorantenkriterium (b) gegen eine geometrische Reihe (vgl. 16.4 (a)) abgeschätzt wird.

□

16.8 Bemerkung:

(a) Um nach Quotienten- bzw. Wurzelkriterium absolute Konvergenz von $\sum_{k=1}^{\infty} a_k$ folgern zu können, reicht es aus zu zeigen, dass

$$\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| < 1 \quad \text{bzw.} \quad \lim_{k \rightarrow \infty} \sqrt[k]{|a_k|} < 1 \text{ ist.}$$

(b) Gilt

$$\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| > 1 \quad \text{oder} \quad \lim_{k \rightarrow \infty} \sqrt[k]{|a_k|} > 1,$$

so ist $\sum_{k=1}^{\infty} a_k$ divergent.

16.9 Beispiel:

(a) Für jede reelle Zahl z ist $\sum_{k=0}^{\infty} \underbrace{\frac{z^k}{k!}}_{=: a_k} \left(= 1 + \sum_{k=1}^{\infty} \frac{z^k}{k!} \right)$ absolut konvergent. Es

ist nämlich für $z \neq 0$

$$\begin{aligned} \left| \frac{a_{k+1}}{a_k} \right| &= \left| \frac{z^{k+1} k!}{z^k (k+1)!} \right| \\ &= \frac{|z|}{k+1} \rightarrow 0 \text{ für } k \rightarrow \infty \end{aligned}$$

und damit nach Quotientenkriterium $\sum_{k=0}^{\infty} a_k$ absolut konvergent. (Fall $z = 0$: trivial.)

Anmerkung: Es gilt $\sum_{k=0}^{\infty} \frac{z^k}{k!} = e^z \quad \forall z \in \mathbb{R}$ (e: Euler'sche Zahl, vgl. 14.15)

(b) Wir betrachten $\sum_{k=2}^{\infty} \underbrace{\frac{1}{(\ln k)^k}}_{=: a_k}$. Wegen $\sqrt[k]{|a_k|} = \frac{1}{\ln k} \rightarrow 0$ für $k \rightarrow \infty$ konvergiert die Reihe nach Wurzelkriterium absolut.

(c) $\sum_{k=1}^{\infty} \frac{1}{k(k+1)}$ ist absolut konvergent, und $\sum_{k=1}^n \frac{1}{k(k+1)} = 1$, denn:

$$\sum_{k=1}^n \frac{1}{k(k+1)} = \sum_{k=1}^n \left(\frac{1}{k} - \frac{1}{k+1} \right) = 1 - \frac{1}{n+1}.$$

(d) Für jedes $r \in \mathbb{N}$, $r \geq 2$, ist $\sum_{k=1}^{\infty} \frac{1}{k^r}$ absolut konvergent, denn

$$\begin{aligned} \sum_{k=1}^n \frac{1}{k^r} &\leq \sum_{k=1}^n \frac{1}{k^2} \\ &< 1 + \sum_{k=2}^n \frac{1}{(k-1)k} < 2. \quad (\text{vgl. (a)}) \end{aligned}$$

Anmerkung: 1: Man kann sogar zeigen, dass für jedes $r \in \mathbb{R}, r > 1$, $\sum_{k=1}^n \frac{1}{k^r}$ absolut konvergiert.

Anmerkung: 2: $\sum_{k=1}^{\infty} \frac{1}{k^2} = \frac{\pi^2}{6}$, $\sum_{k=1}^{\infty} \frac{1}{k^4} = \frac{\pi^4}{90}$, $\sum_{k=1}^{\infty} \frac{1}{k^6} = \frac{\pi^6}{945}$.

Anmerkung: 3: Wegen $\lim_{k \rightarrow \infty} \left| \frac{k^r}{(k+1)^r} \right| = 1$ nützt das Quotientenkriterium hier nichts.

(e) **Warnung:** Für $\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| = 1$ bzw. $\lim_{k \rightarrow \infty} \sqrt[k]{|a_k|} = 1$ ist kein Schluss auf das Konvergenzverhalten möglich.

Für die alternierende harmonische Reihe $\sum_{k=1}^{\infty} \frac{(-1)^k}{k}$ (vgl. 16.4 (c)) gilt z.B.

$\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| = \lim_{k \rightarrow \infty} \left| \frac{k}{k+1} \right| = 1$. Die Reihe konvergiert, jedoch nicht absolut.

Ebenso gilt

$$\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| = 1 \text{ für } \sum_{k=1}^{\infty} \frac{1}{k^r}, r > 1$$

(absolut konvergent, siehe (b)) oder $0 < r < 1$ (divergiert).

16.10 Satz: (Umordnungssatz)

Ist die Reihe $\sum_{k=1}^{\infty} a_k$ absolut konvergent, und ist $\sigma : \mathbb{N} \rightarrow \mathbb{N}$ eine beliebige Permutation (bijektive Abbildung, „Umnummerierung“), so ist auch $\sum_{k=1}^{\infty} a_{\sigma_k}$ absolut konvergent, und es gilt

$$\sum_{k=1}^{\infty} a_k = \sum_{k=1}^{\infty} a_{\sigma_k}$$

Beweis:

Zu jedem $m \in \mathbb{N}$ gibt es ein $N \in \mathbb{N}$, so dass $\{\sigma_1, \sigma_2, \dots, \sigma_m\} \subset \{1, \dots, N\}$ gilt und damit:

$$\begin{aligned} \sum_{k=1}^m |a_{\sigma_k}| &\leq \sum_{k=1}^N |a_k| \\ &\leq \sum_{k=1}^{\infty} |a_k| =: S \end{aligned}$$

Nach Satz 16.7 (a) ist $\sum_{k=1}^m |a_{\sigma_k}|$ absolut konvergent; für den Grenzwert $S' := \sum_{k=1}^{\infty} |a_{\sigma_k}|$ gilt $S' \leq S$.

Umgekehrt ist auch $\sum_{k=1}^{\infty} |a_k|$ eine Umordnung von $\sum_{k=1}^{\infty} |a_{\sigma_k}|$, also $S \leq S' \Rightarrow S = S'$.

Wir betrachten nun die absolut konvergente Reihe $\sum_{k=1}^{\infty} (|a_k| + a_k)$ und erhalten

$$\sum_{k=1}^{\infty} (|a_k| + a_k) = \sum_{k=1}^{\infty} (|a_{\sigma_k}| + a_{\sigma_k}).$$

Wegen der Linearität (vgl. 16.3 (c)) folgt

$$S + \sum_{k=1}^{\infty} a_k = S' + \sum_{k=1}^{\infty} a_{\sigma_k}$$

und daher

$$\sum_{k=1}^{\infty} a_k = \sum_{k=1}^{\infty} a_{\sigma_k}.$$

□

16.11 Bemerkung:

Wenn für jede Permutation $\sigma : \mathbb{N} \rightarrow \mathbb{N}$ die Reihe $\sum_{k=1}^{\infty} a_{\sigma_k}$ konvergiert, so konvergiert die Reihe $\sum_{k=1}^{\infty} a_k$ absolut.

16.12 Umordnung nicht absolut konvergenter Reihen

Ordnet man eine nicht absolut konvergente Reihe um, so kann sich der Grenzwert und das Konvergenzverhalten ändern.

Beispiel: Eine Umordnung der alternierenden harmonischen Reihe:

$$\begin{aligned}
 & 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \frac{1}{5} - \frac{1}{6} \pm \dots \\
 & \rightarrow 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} \\
 & \quad + \frac{1}{5} + \frac{1}{7} - \frac{1}{6} \\
 & \quad + \frac{1}{9} + \frac{1}{11} + \frac{1}{13} + \frac{1}{15} - \frac{1}{8} \\
 & \quad + \frac{1}{17} + \frac{1}{19} + \frac{1}{21} + \frac{1}{23} + \frac{1}{25} + \frac{1}{27} + \frac{1}{29} + \frac{1}{31} - \frac{1}{10} \\
 & \quad \vdots \\
 & \quad + \frac{1}{2^n + 1} + \frac{1}{2^n + 3} + \frac{1}{2^n + 5} + \frac{1}{2^n + 7} + \dots + \frac{1}{2^{n+1} - 1} - \frac{1}{2n + 2} \\
 & \quad \vdots
 \end{aligned}$$

Da die positiven Teilsummen

$$\frac{1}{2^n + 1} + \frac{1}{2^n + 3} + \dots + \frac{1}{2^{n+1} - 1} > 2^{n-1} \cdot \frac{1}{2^{n+1}} = \frac{1}{4}$$

erfüllen, ist diese Reihe nach dem Cauchy-Kriterium 14.11 (b) divergent.

Es gilt sogar der Satz:

Ist eine Reihe konvergent, jedoch nicht absolut konvergent, dann gibt es

- zu jeder reellen Zahl g eine Umordnung mit g als Grenzwert,
- bestimmt divergente Umordnungen mit Grenzwert $+\infty$ und $-\infty$,
- eine unbestimmt divergente Umordnung.

Jetzt wollen wir uns noch mit der „Multiplikation“ von Reihen befassen. Kann man das Produkt zweier Reihen „ausmultiplizieren“, etwa folgendermaßen

$$\left(\sum_{k=1}^{\infty} a_k \right) \left(\sum_{l=1}^{\infty} b_l \right) = \sum_{k=1}^{\infty} a_k \cdot b_l \quad ?$$

(Dabei soll die Reihe auf der rechten Seite jedes Paar (k, l) , $k, l \in \mathbb{N}$ von Indizes genau einmal enthalten).

Die Antwort gibt der folgende Satz:

16.13 Satz: (Produkt von Reihen)

Die Reihen $\sum_{k=1}^{\infty} a_k$, $\sum_{l=1}^{\infty} b_l$ seien absolut konvergent.

Dann ist für jede bijektive Abbildung $(\sigma, \mu) : \mathbb{N} \rightarrow \mathbb{N}^2$ der Indexpaare die Reihe $\sum_{k=1}^{\infty} a_{\sigma_k} b_{\mu_k}$ absolut konvergent, und es gilt

$$\sum_{k=1}^{\infty} a_{\sigma_k} b_{\mu_k} = \left(\sum_{k=1}^{\infty} a_k \right) \left(\sum_{l=1}^{\infty} b_l \right).$$

(ohne Beweis)

Bemerkung: Für nicht absolut konvergente Reihen ist aufgrund des Satzes aus 16.12 auf Seite 131 keine vergleichbare Aussage möglich.

16.14 Folgerung:

Mit der Reihenfolge der Indexpaare gemäß

σ_k	0	1	2	3	...
μ_k					
0	0	2	5	9	
1	1	4	8		
2	3	7			
3	6				
$\vdots$					

erhält man (hier wird ab 0 nummeriert, damit sich einfachere Ausdrücke ergeben!)

$$\begin{aligned} \left(\sum_{k=0}^{\infty} a_k \right) \left(\sum_{l=0}^{\infty} b_l \right) &= \sum_{n=0}^{\infty} \left(\sum_{k=0}^n a_k b_{n-k} \right) \\ &= a_0 b_0 + (a_0 b_1 + a_1 b_0) + (a_0 b_2 + a_1 b_1 + a_2 b_0) + \dots \end{aligned}$$

16.15 Beispiel:

Nach 16.4 (a) gilt für $z \in \mathbb{R}, 0 < z < 1$,

$$\sum_{k=0}^{\infty} z^k = \frac{1}{1-z} \quad \text{und} \quad \sum_{l=0}^{\infty} (-z)^l = \frac{1}{1+z}$$

Für das Produkt erhalten wir nach 16.14

$$\begin{aligned}\left(\sum_{k=0}^{\infty} z^k\right) \left(\sum_{l=0}^{\infty} (-z)^l\right) &= \sum_{n=0}^{\infty} \left(\sum_{k=0}^n z^k (-z)^{n-k}\right) \\ &= \sum_{n=0}^{\infty} \left(\sum_{k=0}^n (-1)^{n-k} z^n\right)\end{aligned}$$

und wegen $\sum_{k=0}^n (-1)^{n-k} = \begin{cases} 1 & \text{für } n \text{ gerade} \\ 0 & \text{für } n \text{ ungerade} \end{cases}$

$$\left(\sum_{k=0}^{\infty} z^k\right) \left(\sum_{l=0}^{\infty} (-z)^l\right) = \sum_{n=0}^{\infty} z^{2n}$$

Nach 16.4 (a) ist $\sum_{k=0}^{\infty} z^{2k} = \frac{1}{1-z^2}$.

Andererseits gilt nach der dritten binomischen Formel

$$\frac{1}{1-z} \cdot \frac{1}{1+z} = \frac{1}{1-z^2}.$$

Kapitel 17

Potenzreihen

17.1 Motivation

Potenzreihen sind wichtig bei der Darstellung und Approximation von Funktionen. (Taylorreihen, erzeugende Funktionen).

17.2 Definition: (Potenzreihe, Entwicklungspunkt)

Eine Reihe der Form

$$\sum_{i=0}^{\infty} a_i (x - x_0)^i \quad \text{mit } x_0, x \in \mathbb{R}$$

heißt Potenzreihe. x_0 heißt Entwicklungspunkt. Im folgenden sei stets $x_0 = 0$.

17.3 Beispiel:

(a) Die Exponentialreihe $\sum_{i=0}^{\infty} \frac{x^i}{i!}$ konvergiert für alle $x \in \mathbb{R}$.

(b) Die geometrische Reihe stellt für $|x| < 1$ eine Funktion dar:

$$\sum_{i=1}^{\infty} x^i = \frac{1}{1-x} \quad (\text{vgl. 16.15}).$$

Die Reihe konvergiert absolut für $|x| < 1$ und divergiert für $|x| \geq 1$. Es scheint eine einfache „Grenze“ zu geben, welche die Bereiche trennt, in denen eine Potenzreihe konvergiert oder divergiert.

17.4 Satz:

Für eine Potenzreihe, die nicht für alle $x \in \mathbb{R}$ konvergiert, gibt es ein $R \geq 0$, so dass sie für $|x| < R$ absolut konvergiert und für $|x| > R$ divergiert.

Beweis:

siehe Brill, S. 271.

Die Zahl R heißt Konvergenzradius und das Intervall $] -R, R[$ heißt Konvergenzintervall.

Man legt fest:

- $R = \infty$: Die Potenzreihe konvergiert absolut für alle $x \in \mathbb{R}$.
- $R = 0$: Die Potenzreihe konvergiert nur für $x = 0$.

17.5 Anmerkung:

Man darf auch komplexe Zahlen $z \in \mathbb{C}$ einsetzen. In $\mathbb{C}$ ist die Menge $\{z \in \mathbb{C} \mid |z| < R\}$ ein (offener) Kreis.

17.6 Satz: (Hadamard, Berechnung des Konvergenzradius)

Hat die Potenzreihe $\sum_{i=0}^{\infty} a_i x^i$ einen Konvergenzradius R und konvergiert die Folge $\left(\sqrt[n]{|a_n|} \right)_{n \in \mathbb{N}}$, dann gilt:

$$\lim_{n \rightarrow \infty} \sqrt[n]{|a_n|} = \frac{1}{R}.$$

Die Aussage bleibt richtig für $R = 0$ oder $R = \infty$, wenn man $\frac{1}{R}$ als ∞ oder 0 interpretiert.

Beweis:

siehe Brill, S. 272.

17.7 Beispiel:

- (a) Geometrische Reihe $\sum_{i=1}^{\infty} x^i$: $\lim_{n \rightarrow \infty} \sqrt[n]{|a_n|} = \lim_{n \rightarrow \infty} \sqrt[n]{1} = 1 \Rightarrow R = 1$. Aber Divergenz für $x = \pm R$.

(b) $\sum_{i=1}^{\infty} \frac{x^i}{i}$ hat Konvergenzradius $R = 1$, denn:

$$\begin{aligned}\lim_{n \rightarrow \infty} \sqrt[n]{\frac{1}{n}} &= \lim_{n \rightarrow \infty} \frac{1}{\sqrt[n]{n}} \\ &= \frac{1}{1} = 1 \\ \Rightarrow R &= 1.\end{aligned}$$

Für $x = 1$: Divergenz, harmonische Reihe.

Für $x = -1$: Konvergenz, alternierende harmonische Reihe.

Kapitel 18

Darstellung von Zahlen in Zahlssystemen

18.1 Motivation

Die Wahl von 10 als Basis eines Zahlensystems ist nicht zwingend. In der Informatik sind Zahldarstellungen z.B. im Hexadezimalsystem (16), besonders aber im Dualsystem (2) wichtig.

18.2 Satz: *(Darstellung natürlicher Zahlen in Zahlssystemen)*

Sei $b \in \mathbb{N}, b \geq 2$. Dann lässt sich jedes $n \in \mathbb{N}$ eindeutig darstellen in der Form

$$n = a_m b^m + a_{m-1} b^{m-1} + \dots + a_1 b^1 + a_0 b^0$$

mit $m \in \mathbb{N}_0$ und $a_i \in \{0, \dots, b-1\}$.

Anstelle von $a_m b^m + \dots + a_0 b^0$ schreibt man $(a_m a_{m-1} \dots a_0)_b$

Beweis:

Hartmann, S. 48.

18.3 Definition: *(Ziffern, Basis)*

Die möglichen Koeffizienten $0, \dots, b-1$ nennt man Ziffern des b -adischen Systems zur Basis b . Stets sei $b^0 = 1$.

18.4 Beispiel:

$$b = 2 : 256 = 2 \cdot 10^2 + 5 \cdot 10 + 6 \cdot 1 = 1 \cdot 2^8 + 0 \cdot 2^7 + \dots + 0 \cdot 1 = (100000000)_2$$

Wie findet man die Ziffern für die Zahldarstellung?

Durch wiederholte Division mit Rest:

$$\begin{array}{rclcl}
 (a_m b^m + \dots + a_0 b^0) : b & = & a_m b^{m-1} + \dots + a_1 b^0 & \text{Rest } a_0 \\
 \swarrow & & \nearrow & \\
 (a_m b^{m-1} + \dots + a_1 b^0) : b & = & a_m b^{m-2} + \dots + a_2 b^0 & \text{Rest } a_1 \\
 \swarrow & & \nearrow & \\
 & \vdots & & \\
 a_m b^0 : b & = & 0 & \text{Rest } a_m
 \end{array}$$

18.5 Definition: (*b*-adischer Bruch)

Es sei $z \in \mathbb{Z}$ und für alle $n \geq z$ sei $a_n \in \{0, \dots, b-1\}$. Eine Reihe der Form

$$\sum_{n=z}^{\infty} \frac{a_n}{b^n} \quad (\text{Reihe mit negativen Indizes})$$

mit $a_z \neq 0$ heißt b-adischer Bruch. Man schreibt dafür $(a_z a_{z+1} \dots E - z)_b$. Man spricht von

- einem endlichen *b*-adischen Bruch und schreibt $(a_z a_{z+1} \dots a_n E - z)_b$, falls für alle $i > n$ gilt: $a_i = 0$.
- einem periodischen *b*-adischen Bruch und setzt

$$(a_z, a_{z+1} \dots a_n \overline{p_1 \dots p_r} E - z)_b := (a_z, \dots a_n p_1 \dots p_r p_1 \dots p_r \dots E - z)_b,$$

falls eine Folge $p_1 \dots p_r$ von Ziffern sich ständig wiederholt.

Alternative Schreibweisen:

$$\begin{array}{l}
 (a_z a_{z+1} \dots a_0, a_1 a_2 \dots)_b \quad \text{für } z \leq 0, \\
 \text{und} \\
 (0, \underbrace{0 \dots 0}_{z-1 \text{ Stück}} a_z a_{z+1} \dots)_b \quad \text{für } z > 0.
 \end{array}$$

18.6 Beispiel:

$$\begin{aligned}(4,625E1)_7 &= (46,25)_7 = 4 \cdot 7^1 + 6 \cdot 7^0 + \frac{2}{7} + \frac{5}{7^2} \\ &= \frac{1685}{49}.\end{aligned}$$

18.7 Satz: (Darstellung reeller Zahlen in Zahlensystemen)

Sei $b \in \mathbb{N}, b \geq 2$. Dann gilt:

- (a) Jeder b -adische Bruch konvergiert gegen ein $x \in \mathbb{R}, x \geq 0$.
- (b) Zu jedem $x \in \mathbb{R}, x \geq 0$, gibt es einen b -adischen Bruch, der gegen x konvergiert.
- (c) Jeder periodische b -adische Bruch konvergiert gegen ein $y \in \mathbb{Q}, y \geq 0$.
- (d) Zu jedem $y \in \mathbb{Q}, y \geq 0$ gibt es einen periodischen oder endlichen b -adischen Bruch, der gegen y konvergiert.

Beweis:

(a), (b), (c) siehe Hartmann, S. 251.

Zu (d): o.B.d.A. sei $y = \frac{p}{q} < 1, p \in \mathbb{N}_0, q \in \mathbb{N}$.

Gesucht $a_1, a_2 \dots$ mit

$$\frac{p}{q} = \sum_{n=1}^{\infty} \frac{a_n}{b^n}.$$

Berechnung der a_i durch folgenden Divisionsalgorithmus für natürliche Zahlen:

$$\begin{array}{llll} b \cdot p & = & a_1 \cdot q + r_1 & 0 \leq r_1 < q, \\ b \cdot r_1 & = & a_2 \cdot q + r_2 & 0 \leq r_2 < q \\ b \cdot r_2 & = & a_3 \cdot q + r_3 & 0 \leq r_3 < q \\ b \cdot r_3 & = & a_4 \cdot q + r_4 & 0 \leq r_4 < q \\ & \vdots & & \\ b \cdot r_n & = & a_{n+1} \cdot q + r_{n+1} & 0 \leq r_{n+1} < q. \end{array} \quad (*)$$

Es ist stets $a_i < b$, denn:

Man setze $r_0 := p$. Aus $b \leq a_i$ und $r_{i-1} < q$ für $i \in \mathbb{N}$ folgt

$$b \cdot r_{i-1} < b \cdot q \leq a_i \cdot q$$

$\Rightarrow$ Widerspruch zu

$$b \cdot r_{i-1} = a_i \cdot q + r_i \geq a_i \cdot q.$$

Diese a_i sind die gesuchten Koeffizienten, denn aus dem Gleichungssystem (*) folgt:

$$\begin{array}{rcl}
 p & = & \frac{a_1}{b}q + \frac{r_1}{b} \\
 r_1 & = & \frac{a_2}{b}q + \frac{r_2}{b} \\
 \vdots & & \vdots \\
 r_n & = & \frac{a_{n+1}}{b}q + \frac{r_{n+1}}{b}
 \end{array}
 \quad \begin{array}{c}
 \Rightarrow \\
 r_1 \text{ einsetzen} \\
 \\
 \Rightarrow \\
 r_i \text{ einsetzen}
 \end{array}
 \quad
 \begin{array}{rcl}
 p & = & \left(\frac{a_1}{b} + \frac{a_2}{b^2} \right) q + \frac{r_2}{b^2} \\
 & & \vdots \\
 p & = & \left(\sum_{k=1}^{n+1} \frac{a_k}{b^k} \right) q + \frac{r_{n+1}}{b^{n+1}}
 \end{array}$$

$$\Rightarrow \underbrace{\frac{r_{n+1}}{b^{n+1}}q}_{\substack{\text{konvergiert gegen 0} \\ \text{für } n \rightarrow \infty}} = \frac{p}{q} - \underbrace{\sum_{k=1}^{n+1} \frac{a_k}{b^k}}_{\text{konvergiert gegen } \frac{p}{q}}$$

Wann endet der Algorithmus?

Gilt $r_n = 0$, dann ist $\frac{p}{q} = (0, a_1 \dots a_n)_b$. Gilt stets $r_n \neq 0$, dann muss sich ein Rest wiederholen, der Bruch wird periodisch, alle Divisionen wiederholen sich. Ist z.B. $r_n = r_s$ für ein $n > s$, dann folgt $a_{s+1} = a_{n+1}$, also $\frac{p}{q} = (0, a_1 \dots a_s \overline{a_{s+1} \dots a_n})_b$.

18.8 Bemerkung:

Die Darstellung aus 18.5 auf Seite 140 zur Basis $b=2$ entspricht genau der Speicherung von Zahlen im Computer. Die **endliche (!)** Folge der a_i heißt Mantisse. Die sogenannte Gleitkommadarstellung umfasst das Vorzeichen, die Mantisse und den Exponenten. Typischer **float**:

$$\underbrace{-1}_{1 \text{ Bit}}, \underbrace{01 \dots 1}_{23 \text{ Bit}} \cdot 2^{\underbrace{99}_{8 \text{ Bit}}}$$

Kapitel 19

Binomialkoeffizienten und die Binomialreihe

19.1 Motivation

Binomialkoeffizienten spielen eine große Rolle

- in der Kombinatorik, wenn man die Anzahl der k -elementigen Teilmengen einer n -elementigen Menge sucht
- beim binomischen Satz zur Berechnung von $(a + b)^n$, $n \in \mathbb{N}$.
- in der Potenzreihendarstellung von $(1 + x)^\alpha$ für $|x| < 1$.

19.2 Definition: (*Binomialkoeffizient*)

Seien $n, k \in \mathbb{N}_0$. Dann heißt

$$\binom{n}{k} = \begin{cases} \text{Anzahl der } k\text{-elementigen Teilmengen} \\ \text{einer } n\text{-elementigen Menge} \end{cases}$$

den Binomialkoeffizient n über k .

Bemerkung: Offensichtlich gilt:

$$\begin{aligned} \binom{n}{0} = \binom{n}{n} &= 1, \\ \binom{n}{1} &= n, \\ \binom{n}{k} &= 0 \quad \text{für } k > n \end{aligned}$$

Der folgende Satz liefert eine rekursive Beschreibung der Binomialkoeffizienten.

19.3 Satz: (Rekursionsbeziehung für Binomialkoeffizienten)

Seien $n, k \in \mathbb{N}_0$ und $n \geq k$. Dann gilt:

$$\binom{n}{k} = \begin{cases} 1 & \text{falls } k = 0 \text{ oder } k = n \\ \binom{n-1}{k} + \binom{n-1}{k-1} & \text{sonst} \end{cases}$$

Beweis:

Aussage klar für $k = 0$ oder $k = n$. Sei also $0 < k < n$. Sei $M = \{x_1, \dots, x_n\}$. Die k -elementigen Teilmengen N von M zerfallen in 2 Klassen:

- (a) Mengen N mit $x_n \notin N$. Diese entsprechen den k -elementigen Teilmengen von $M \setminus \{x_n\}$. Davon gibt es nach Def. 19.2 $\binom{n-1}{k}$ Stück.
- (b) Mengen N mit $x_n \in N$. Sie entsprechen den $(k-1)$ -elementigen Teilmengen von $M \setminus \{x_n\}$. Davon gibt es $\binom{n-1}{k-1}$ Stück.

Insgesamt gilt daher:

$$\binom{n}{k} = \binom{n-1}{k} + \binom{n-1}{k-1}.$$

□

19.4 Pascal'sches Dreieck

Satz 19.3 erlaubt die Berechnung von Binomialkoeffizienten mit Hilfe eines Dreiecks, in dem jeder Koeffizient als Summe der beiden schräg darüber stehenden Koeffizienten berechnet wird.

$$\begin{array}{ccccccc}
 & & & & \binom{0}{0} = 1 & & \\
 & & & \binom{1}{0} = 1 & & \binom{1}{1} = 1 & \\
 & & \swarrow & \searrow & \swarrow & \searrow & \\
 & \binom{2}{0} = 1 & & \binom{2}{1} = 2 & & \binom{2}{2} = 1 & \\
 & \swarrow & \searrow & \swarrow & \searrow & \swarrow & \searrow \\
 \binom{3}{0} = 1 & & \binom{3}{1} = 3 & & \binom{3}{2} = 3 & & \binom{3}{3} = 1
 \end{array}$$

Für große n ist diese Rekursionsformel ineffizient. Gibt es eine direkte Formel?

19.5 Satz: (Direkte Formel für Binomialkoeffizienten)

Seien $n, k \in \mathbb{N}_0$, $n \geq k$. Dann gilt:

$$\binom{n}{k} = \frac{n!}{k!(n-k)!}.$$

Dabei bezeichnet $k! := k \cdot (k-1) \cdot (k-2) \cdot \dots \cdot 1$ die Fakultät von k . Man setzt $0! := 1$.

Beweis:

Vollständige Induktion

19.6 Beispiel:

Wie viele Möglichkeiten gibt es im Lotto, 6 aus 49 Zahlen auszuwählen?

$$\binom{49}{6} = \frac{49!}{6! \cdot 43!} = \frac{49 \cdot 48 \cdot 47 \cdot 46 \cdot 45 \cdot 44}{6 \cdot 5 \cdot 4 \cdot 3 \cdot 2 \cdot 1} = 13.983.816$$

Dabei spielt die Reihenfolge der 6 Zahlen keine Rolle.

Was hat das mit Analysis zu tun ?

19.7 Satz: (Binomialsatz)

Seien $a, b \in \mathbb{R}$ und $n \in \mathbb{N}_0$. Dann gilt

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

19.8 Beispiel:

(a)

$$\begin{aligned} (a+b)^2 &= \binom{2}{0} a^0 b^2 + \binom{2}{1} a^1 b^1 + \binom{2}{2} a^2 b^0 \\ &= b^2 + 2ab + a^2 \end{aligned}$$

(b)

$$\begin{aligned} (a+b)^3 &= \binom{3}{0} a^0 b^3 + \binom{3}{1} a^1 b^2 + \binom{3}{2} a^2 b^1 + \binom{3}{3} a^3 b^0 \\ &= b^3 + 3ab^2 + 3a^2b + a^3 \end{aligned}$$

19.9 Beweis des Binomialsatzes

Wir verwenden das Prinzip der Indexverschiebung.

$$\sum_{k=0}^n a_k = a_0 + a_1 + \dots + a_k = \sum_{k=1}^{n+1} a_{k-1} \quad (*)$$

um den Binomialsatz mit vollständiger Induktion zu beweisen.

Induktionsanfang:

$$\begin{aligned} \sum_{k=0}^0 \binom{n}{k} a^k b^{n-k} &= \binom{0}{0} a^0 b^0 = 1 \cdot 1 \cdot 1 = 1 \\ &= (a+b)^0. \end{aligned}$$

Induktionsannahme: Es sei

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

Induktionsschluss:

$$\begin{aligned} (a+b)^{n+1} &= (a+b)^n (a+b) \\ &\stackrel{\text{Ind. Ann.}}{=} \left(\sum_{k=0}^n \binom{n}{k} a^k b^{n-k} \right) (a+b) \\ &= \sum_{k=0}^n \binom{n}{k} a^{k+1} b^{n-k} + \sum_{k=0}^n \binom{n}{k} a^k b^{n-k+1} \\ &\stackrel{(*)}{=} \sum_{k=0}^{n+1} \binom{n}{k-1} a^k b^{n-k+1} + \sum_{k=0}^n \binom{n}{k} a^k b^{n-k+1} \\ &\stackrel{\text{abspalten}}{=} \sum_{k=0}^n \left[\binom{n}{k-1} + \binom{n}{k} \right] a^k b^{n-k+1} + \\ &\quad \underbrace{\binom{n}{n}}_1 a^{n+1} b^0 + \underbrace{\binom{n}{0}}_1 a^0 b^{n+1} \\ &\stackrel{19.3}{=} \sum_{k=0}^n \binom{n+1}{k} a^k b^{n-k+1} + \\ &\quad \binom{n+1}{n+1} a^{n+1} b^0 + \binom{n+1}{0} a^0 b^{n+1} \\ &= \sum_{k=0}^{n+1} \binom{n+1}{k} a^k b^{n+1-k}. \end{aligned}$$

□

19.10 Binomialreihe: Motivation

Aus dem Binomialsatz folgt für $n \in \mathbb{N}_0$:

$$(1+x)^n = (x+1)^n = \sum_{k=0}^n \binom{n}{k} x^k 1^{n-k} = \sum_{k=0}^n \binom{n}{k} x^k.$$

Kann man eine ähnliche Beziehung auch für $(1+x)^\alpha$ gewinnen, wenn α keine natürliche Zahl ist? Hierzu müssen wir den Binomialkoeffizienten

$$\binom{n}{k} = \frac{n!}{k!(n-k)!} = \frac{1}{k!} \prod_{j=0}^{k-1} (n-j)$$

verallgemeinern zu

$$\boxed{\binom{\alpha}{k} := \frac{1}{k!} \prod_{j=0}^{k-1} (\alpha - j)} \quad \alpha \in \mathbb{R}, k \geq 0.$$

Dann kann man zeigen

19.11 Satz: (Binomialreihe)

Sei $\alpha \in \mathbb{R}$ und $-1 < x < 1$. Dann hat $(1+x)^\alpha$ die Potenzreihenentwicklung

$$(1+x)^\alpha = \sum_{k=0}^{\infty} \binom{\alpha}{k} x^k$$

Bemerkung: Für $\alpha \in \mathbb{N}_0$ bricht die Reihe wegen $\binom{\alpha}{n} = 0 \quad \forall \alpha > n$ ab (vgl. 19.2 auf Seite 143).

19.12 Konvergenz der Binomialreihe

Sei $\alpha \notin \mathbb{N}_0$ und $x \neq 0$. Dann gilt mit dem Quotientenkriterium für $a_k := \binom{\alpha}{k} x^k$:

$$\begin{aligned} \left| \frac{a_{k+1}}{a_k} \right| &= \left| \frac{\binom{\alpha}{k+1} x^{k+1}}{\binom{\alpha}{k} x^k} \right| \\ &= \left| \frac{\frac{1}{(k+1)!} \prod_{j=0}^k (\alpha - j)}{\frac{1}{k!} \prod_{j=0}^{k-1} (\alpha - j)} \right| \cdot |x| \\ &= \left| \frac{\alpha - k}{k+1} \right| \cdot |x|. \end{aligned}$$

Wegen $\lim_{k \rightarrow \infty} \left| \frac{\alpha - k}{k+1} \right| = |-1| = 1$ ist $\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| = |x|$.

Damit konvergiert (a_k) nach dem Quotientenkriterium für $|x| < 1$.

19.13 Beispiele

(a) $\alpha = -1$:

$$\begin{aligned} \binom{-1}{k} &= \frac{1}{k!} \prod_{j=0}^{k-1} (-1-j) = \frac{1}{1 \cdot 2 \cdot \dots \cdot k} (-1) \cdot (-2) \cdot \dots \cdot (-k) = (-1)^k \\ \Rightarrow \frac{1}{1+x} &= \sum_{k=0}^{\infty} (-1)^k x^k \\ &= \sum_{k=0}^{\infty} (-x)^k \quad \text{für } |x| < 1. \end{aligned}$$

Die geometrische Reihe ist also ein Spezialfall der Binomialreihe.

(b) $\alpha = \frac{1}{2}$:

$$\begin{aligned} \binom{\frac{1}{2}}{0} &= \frac{1}{0!} \underbrace{\prod_{j=0}^{-1} \left(\frac{1}{2} - j\right)}_1 = 1 \\ \binom{\frac{1}{2}}{1} &= \frac{1}{1!} \prod_{j=0}^0 \left(\frac{1}{2} - j\right) = \frac{1}{2} \\ \binom{\frac{1}{2}}{2} &= \frac{1}{2!} \prod_{j=0}^1 \left(\frac{1}{2} - j\right) = \frac{1}{2} \cdot \frac{1}{2} \cdot \left(-\frac{1}{2}\right) = -\frac{1}{8} \\ \Rightarrow \sqrt{1+x} &= 1 + \frac{1}{2}x - \frac{1}{8}x^2 + \mathcal{O}(x^3) \quad \text{für } |x| < 1. \end{aligned}$$

Derartige Potenzreihendarstellungen sind nützlich zur numerischen Approximation der Wurzelfunktion.

Kapitel 20

Stetigkeit

20.1 Motivation

- In technischen Systemen erwartet man häufig, dass sich das Resultat nur wenig ändert, wenn man die Eingabegrößen nur gering variiert.
- Mathematisch kann man dies durch das Konzept der Stetigkeit formalisieren.

20.2 Definition: (Konvergenz gegen Grenzwert)

Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ eine Funktion und $\xi \in \mathbb{R}$. Wir sagen, $f(x)$ konvergiert für $x \rightarrow \xi$ gegen den Grenzwert η , falls für jede (!) Folge $(x_n)_{n \in \mathbb{N}}$ mit $x_n \neq \xi \forall n \in \mathbb{N}$ gilt:

$$\lim_{n \rightarrow \infty} x_n = \xi \quad \Rightarrow \quad \lim_{n \rightarrow \infty} f(x_n) = \eta.$$

In diesem Falls schreiben wir $\lim_{x_n \rightarrow \xi} f(x_n) = \eta$.

20.3 Beispiele

- (a) $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$ hat für jedes $\xi \in \mathbb{R}$ einen Grenzwert:
Sei (x_n) eine Folge mit $\lim_{n \rightarrow \infty} x_n = \xi$.

Dann konvergiert die Folge (x_n^2) gegen ξ^2 .

- (b) Die Sprungfunktion

$$f(x) := \begin{cases} 0 & \text{für } x < 0 \\ 1 & \text{für } x \geq 0 \end{cases}$$

hat in 0 keinen Grenzwert.

Sei (x_n) eine Folge mit $x_n < 0 \forall n \in \mathbb{N}$ und $\lim_{n \rightarrow \infty} x_n = 0$

$$\Rightarrow \lim_{n \rightarrow \infty} f(x_n) = 0.$$

Für eine Folge x_n mit $x_n \geq 0$ und $\lim_{n \rightarrow \infty} x_n = 0$ gilt jedoch

$$\lim_{n \rightarrow \infty} f(x_n) = 1.$$

Die für Folgen bekannten Grenzwertsätze lassen sich auf Grenzwerte von Funktionen übertragen:

20.4 Satz: (Grenzwertsätze für Funktionen)

Existieren für die Funktionen $f, g : \mathbb{R} \rightarrow \mathbb{R}$ die Grenzwerte im Punkt $\xi \in \mathbb{R}$, so gilt:

$$(a) \quad \lim_{x \rightarrow \xi} (f(x) \pm g(x)) = f(\xi) \pm g(\xi)$$

$$(b) \quad \lim_{x \rightarrow \xi} (f(x) \cdot g(x)) = f(\xi) \cdot g(\xi)$$

$$(c) \quad \lim_{x \rightarrow \xi} \frac{f(x)}{g(x)} = \frac{f(\xi)}{g(\xi)}, \text{ falls } g(\xi) \neq 0.$$

$$(d) \quad \lim_{x \rightarrow \xi} (c \cdot f(x)) = c \cdot f(\xi)$$

$$(e) \quad \lim_{x \rightarrow \xi} |x| = |\xi|.$$

20.5 Definition: (Stetigkeit)

Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ heißt stetig in $\xi \in \mathbb{R}$, wenn dort Funktionswert und Grenzwert übereinstimmen:

$$\lim_{x \rightarrow \xi} f(x) = f\left(\lim_{x \rightarrow \xi} x\right) = f(\xi)$$

f heißt stetig, wenn f in allen Punkten stetig ist.

Der folgende Satz zeigt, dass kleine Änderungen im Argument einer stetigen Funktion nur zu kleinen Änderungen der Funktionswerte führen.

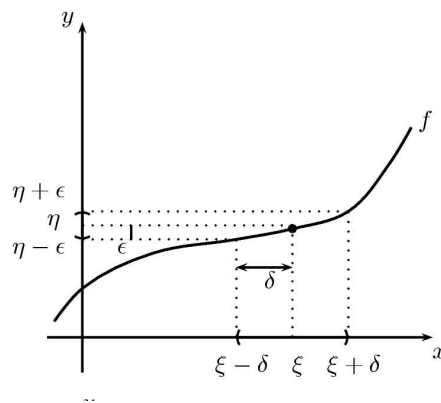
20.6 Satz: ($\epsilon - \delta$ -Kriterium der Stetigkeit)

Für eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ sind äquivalent:

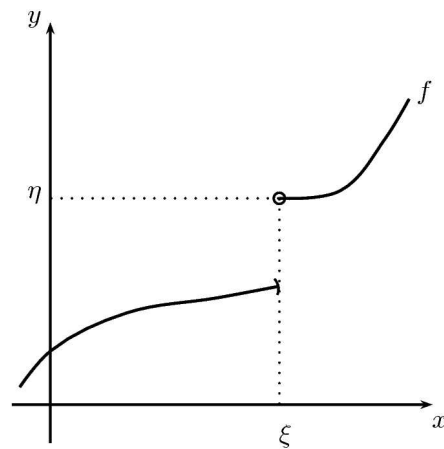
(a) f ist stetig in $\xi \in \mathbb{R}$, d.h. $f(\xi) = \lim_{x \rightarrow \xi} f(x)$.

(b) Zu jedem $\epsilon > 0$ existiert ein $\delta(\epsilon) > 0$ mit

$$|x - \xi| < \delta \Rightarrow |f(x) - f(\xi)| < \epsilon.$$



Man kann also erreichen, dass die Funktionswerte beliebig nahe bei einander liegen, wenn sich die Argumente nur hinreichend wenig unterscheiden.



Das ist hier nicht möglich. Die Funktion ist unstetig in ξ .

Informell:

„Stetige Funktionen kann man zeichnen, ohne abzusetzen“.

20.7 Bemerkung:

- (a) Satz 20.4 auf der vorherigen Seite besagt, dass für in ξ stetige Funktionen $f(x), g(x)$ auch die Funktionen

$f(x) \pm g(x)$, $f(x) \cdot g(x)$, $\frac{f(x)}{g(x)}$ (falls $g(x) \neq 0$) und $c \cdot f(x)$ stetig sind.

Ebenso ist die Betragsfunktion $f(x) = |x|$ stetig.

(b) Die Komposition stetiger Funktionen ist ebenfalls stetig. Denn:

Seien $f, g : \mathbb{R} \rightarrow \mathbb{R}$ stetige Funktionen und (x_n) eine Folge mit $\lim_{n \rightarrow \infty} x_n = \xi$.

$$\Rightarrow \lim_{x_n \rightarrow \xi} f(x_n) = f(\xi), \text{ da } f \text{ stetig.}$$

$$\Rightarrow \lim_{x_n \rightarrow \xi} g(f(x_n)) = g(f(\xi)), \text{ da } g \text{ stetig.}$$

(c) Die Grenzwert- und Stetigkeitsbegriffe sind sinngemäß auch auf Funktionen $f : D \rightarrow \mathbb{R}$ übertragbar, deren Definitionsbereich D eine echte Teilmenge von $\mathbb{R}$ ist. In diesem Fall müssen die betrachteten Folgen (x_n) in D liegen.

20.8 Beispiel:

$f(x) = 5x^3 - 7x + 12$ ist stetig, denn

$$g_1(x) := 12 \text{ ist stetig}$$

$$g_2(x) := -7x \text{ ist stetig}$$

$$g_3(x) := 5x^3 \text{ ist stetig, da Produkt stetiger Funktionen.}$$

$$\Rightarrow f(x) = g_3(x) + g_2(x) + g_1(x) \text{ ist stetig.}$$

20.9 Satz: (Eigenschaften stetiger Funktionen)

Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig. Dann gilt:

(a) **Nullstellensatz:**

Ist $f(a) \cdot f(b) < 0$, so existiert $\xi \in (a, b)$ mit $f(\xi) = 0$.

(b) **Zwischenwertsatz:**

Zu jedem c mit $f(a) < c < f(b)$ existiert $\xi \in (a, b)$ mit $f(\xi) = c$.

(c) **Stetigkeit der Umkehrfunktion:**

Ist $f(x)$ stetig und streng monoton wachsend auf $[a, b]$ (d.h. $x < y \Rightarrow f(x) < f(y)$), so ist auch die Umkehrfunktion

$$f^{-1} : [f(a), f(b)] \rightarrow \mathbb{R}$$

stetig und streng monoton wachsend.

(d) **Maximum-Minimum-Eigenschaft:**

Es existieren $x_1, x_2 \in [a, b]$ mit

$$\begin{aligned}f(x_1) &= \max_{x \in [a, b]} f(x) \\f(x_2) &= \min_{x \in [a, b]} f(x).\end{aligned}$$

Beweis:

von (a): Sei o.B.d.A. $f(a) < f(b)$.

Wir konstruieren induktiv eine Intervallschachtelung (Folge von Intervallen $[a_n, b_n]$) mit folgenden Eigenschaften:

- (a) $f(a_n) \leq 0, f(b_n) \geq 0 \quad \forall n$
- (b) $a \leq a_{n-1} \leq a_n < b_n \leq b_{n-1} \leq b$
- (c) $b_n - a_n = \frac{1}{2^n}(b - a)$

Initialisierung:

$$a_0 := a, \quad b_0 := b.$$

Annahme:

Die Intervallschachtelung sei bis zum Index n konstruiert und erfülle (a) - (c).

Sei $c := \frac{a_n + b_n}{2}$.

Ist $f(c) < 0$, setze $a_{n+1} := c, b_{n+1} := b_n$

Ist $f(c) \geq 0$, setze $a_{n+1} := a_n, b_{n+1} := c$.

$\Rightarrow$

- (a) $f(a_{n+1}) \leq 0, f(b_{n+1}) \geq 0$.
- (b) $a_n \leq a_{n+1} < b_{n+1} \leq b_n$.
- (c) $b_{n+1} - a_{n+1} = \frac{1}{2}(b_n - a_n) \stackrel{Ind. Ann.}{=} \frac{1}{2^{n+1}}(b - a)$.

Nach Konstruktion gilt:

- (a_n) mon. wachsend, nach oben beschränkt durch b und $f(a_n) \leq 0$
 $\Rightarrow \exists \tilde{a}$ mit $\lim_{n \rightarrow \infty} a_n = \tilde{a}$ und $f(\tilde{a}) \leq 0$.
- (b_n) mon. fallend, nach unten beschränkt durch a und $f(b_n) \geq 0$
 $\Rightarrow \exists \tilde{b}$ mit $\lim_{n \rightarrow \infty} b_n = \tilde{b}$ und $f(\tilde{b}) \geq 0$.

Wegen $\lim_{n \rightarrow \infty} (b_n - a_n) = \lim_{n \rightarrow \infty} \frac{1}{2^n} (b - a) = 0$ ist $\tilde{a} = \tilde{b} =: \xi$.

Wegen $f(\xi) = f(\tilde{a}) \leq 0$ und $f(\xi) = f(\tilde{b}) \geq 0$ ist $f(\xi) = 0$.

□

Bemerkung: Dieser konstruktive Beweis beschreibt ein numerisches Verfahren zur Nullstellensuche, das so genannte Bisektionsverfahren (Übungsaufgabe).

20.10 Gleichmäßige Stetigkeit

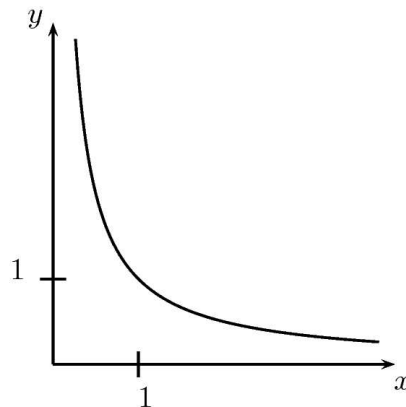
Nach dem $\epsilon - \delta$ -Kriterium (20.6 auf Seite 150) sind stetige Funktionen solche, deren Funktionswert sich bei hinreichend kleinen Änderungen des Arguments nur beliebig wenig ändert; allerdings kann zu vorgegebenen $\epsilon > 0$ (Funktionswertänderung) das zugehörige $\delta(\epsilon)$ (Argumentänderung) noch von ξ abhängen. In manchen Zusammenhängen ist es aber wichtig, dass δ nur von ϵ und nicht von ξ abhängt, d.h. dass man in einem ganzen Intervall zu einem ϵ dasselbe $\delta(\epsilon)$ wählen kann.

Definition:

Gleichmäßige Stetigkeit Eine Funktion $f : D \rightarrow \mathbb{R}$ ($D \subset \mathbb{R}$ Definitionsbereich) heißt gleichmäßig stetig auf D , wenn zu jedem $\epsilon > 0$ ein $\delta(\epsilon) > 0$ existiert, so dass

$$|x_1 - x_2| < \delta \Rightarrow |f(x_1) - f(x_2)| < \epsilon \quad \text{für } x_1, x_2 \in D.$$

Beispiel: $D = (0, +\infty)$, $y = f(x) = \frac{1}{x}$



Zu jedem $\delta > 0$ gibt es $n \in \mathbb{N}$ mit $\frac{1}{n} - \frac{1}{n+1} = \frac{1}{n(n+1)} < \delta$; es ist aber

$$\left| f\left(\frac{1}{n}\right) - f\left(\frac{1}{n+1}\right) \right| = |n - (n+1)| = 1.$$

Damit gibt es zu $0 < \epsilon \leq 1$ kein $\delta(\epsilon)$ mit der geforderten Eigenschaft; f ist also nicht gleichmäßig stetig auf D .

Satz:

Jede stetige Funktion $f : [a, b] \rightarrow \mathbb{R}$ ist gleichmäßig stetig auf $[a, b]$.

Bemerkung: Im obigen Beispiel war der Definitionsbereich kein abgeschlossenes Intervall.

Kapitel 21

Wichtige stetige Funktionen

21.1 Motivation

- Es gibt einige Funktionen, die so wichtig sind, dass sie einen eigenen Namen tragen.
- Wir wollen diese Funktionen und ihre Umkehrfunktionen genauer betrachten.

21.2 Potenzfunktionen

Potenzfunktionen besitzen die Struktur $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^n$ für ein $n \in \mathbb{N}_0$.

Beispiel:

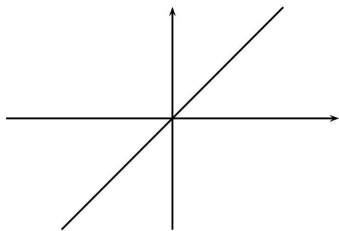


Abbildung 21.1: $n = 1$

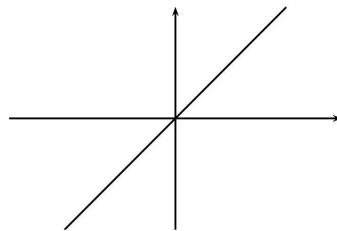


Abbildung 21.2: $n = 2$

Für $n = 0$ definiert man $x^0 := 1$. Wie alle Polynomfunktionen sind Potenzfunktionen stetig. Es gelten die Rechenregeln:

$$(x \cdot y)^n = x^n \cdot y^n, \quad x^n \cdot x^m = x^{n+m}, \quad (x^n)^m = x^{n \cdot m}.$$

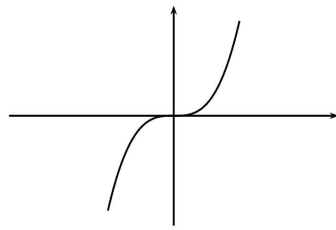


Abbildung 21.3: $n = 3$

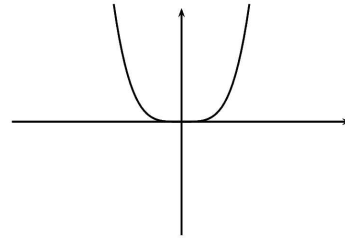


Abbildung 21.4: $n = 4$

21.3 Wurzelfunktionen

Schränkt man die Potenzfunktion $f(x) = x^n$ mit $n \in \mathbb{N}$ auf $\mathbb{R}_0^+$ ein, so ist sie dort nicht nur stetig, sondern auch streng monoton wachsend. Nach 20.9 (d) existiert daher eine stetige, streng monoton wachsende Umkehrfunktion. Die Umkehrfunktion zur n -ten Potenzfunktion

$$f : \mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+, \quad x \mapsto x^n$$

nennt man die n -te Wurzelfunktion

$$g : \mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+, \quad x \mapsto \sqrt[n]{x}$$

Statt $\sqrt[2]{x}$ schreibt man auch $\sqrt{x}$

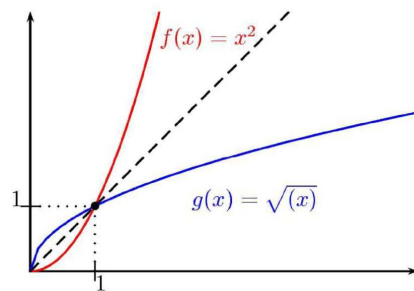


Abbildung 21.5: x^2 und $\sqrt{x}$

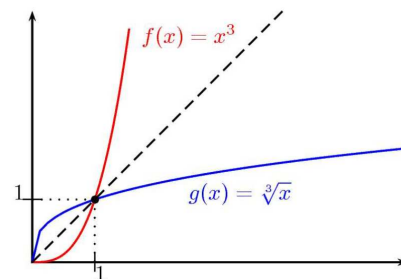


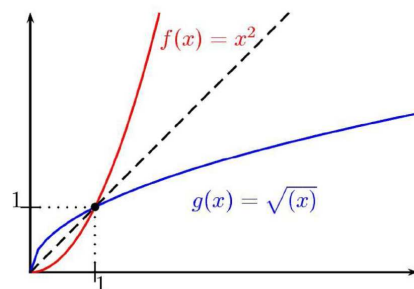
Abbildung 21.6: x^3 und $\sqrt[3]{x}$

Setzt man $x^{\frac{1}{n}} := \sqrt[n]{x}$, $x^{\frac{n}{m}} := \sqrt[m]{x^n}$, $x^{-\frac{n}{m}} := \frac{1}{x^{\frac{n}{m}}}$, so gelten die gleichen Rechenregeln wie für die Potenzfunktionen in 21.2.

21.4 Exponentialfunktion

Sei $a > 0$. Dann bezeichnet man mit $f : \mathbb{R} \rightarrow \mathbb{R}$,

$$f(x) = a^x$$



die Exponentialfunktion zur Basis a (meist ist $a > 1$). Exponentialfunktionen sind stetig.

Es gilt die Funktionalgleichung

$$a^{x+y} = a^x \cdot a^y. \quad (*)$$

Besonders wichtig ist die Exponentialfunktion, bei der die Euler'sche Zahl

$$e = \lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n \approx 2,7182\dots$$

als Basis dient (vgl. 14.15 auf Seite 116). Man bezeichnet sie als

$$\exp(x) := e^x.$$

Für alle $z \in \mathbb{C}$ gilt die Potenzreihendarstellung der Exponentialfunktion:

$$\exp(z) = \sum_{k=0}^{\infty} \frac{z^k}{k!} = 1 + z + \frac{z^2}{2} + \frac{z^3}{6} + \dots$$

Sie konvergiert absolut für alle $z \in \mathbb{C}$ und wächst schneller als jede Potenz:

$$\lim_{x \rightarrow \infty} \frac{e^x}{x^n} = \infty \quad \forall n \in \mathbb{N}.$$

Sei $z = x + iy \in \mathbb{C}$ mit $x, y \in \mathbb{R}$. Dann gilt:

(a) $e^z := e^x \cdot e^{iy}$. Denn $e^z := e^{x+iy} \stackrel{(*)}{=} e^x \cdot e^{iy}$.

(b) $e^x > 0$. Denn:

$$\text{Für } x \geq 0 \quad \text{ist} \quad e^x = \sum_{k=0}^{\infty} \underbrace{\frac{x^k}{k!}}_{\geq 0} \geq 0.$$

$$\text{Für } x < 0 \quad \text{ist} \quad e^x = \frac{1}{\underbrace{e^{-x}}_{\geq 0}} \geq 0.$$

(c) $|e^{iy}| = 1$. Denn: $|e^{iy}|^2 = e^{iy} \overline{e^{iy}} = e^{iy} e^{-iy} = e^{iy-iy} = e^0 = 1$.

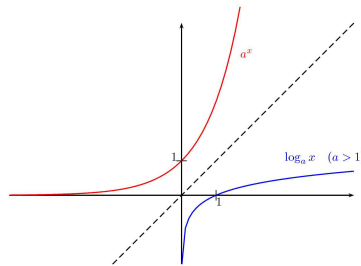
(d) $|e^z| = e^x$. Denn: $|e^z| = |e^x e^{iy}| = \underbrace{|e^x|}_1 \underbrace{|e^{iy}|}_{\geq 0} \stackrel{(c)}{=} |e^x| \stackrel{(b)}{=} e^x$.

21.5 Logarithmusfunktionen

Die Exponentialfunktion $f(x) = a^x$ mit $a > 1$ ist stetig, streng monoton wachsend und bildet $\mathbb{R}$ bijektiv auf $\mathbb{R}^+$ ab. Ihre Umkehrfunktion

$$\log_a : \mathbb{R}^+ \rightarrow \mathbb{R}$$

bezeichnet man als Logarithmusfunktion zur Basis a .



Für $a = e$ ergibt sich der natürliche Logarithmus $\ln$.

Es gelten die Rechenregeln:

$$\begin{aligned}\log_a(x \cdot y) &= \log_a x + \log_a y, \\ \log_a(x^p) &= p \cdot \log_a x.\end{aligned}$$

Logarithmusfunktionen wachsen langsamer als Potenzfunktionen:

$$\lim_{x \rightarrow \infty} \frac{\log_a x}{x^n} = 0 \quad \text{für } a > 1, n \in \mathbb{N}.$$

21.6 Trigonometrische Funktionen

Die Koordinaten eines Punktes auf dem Einheitskreis werden in Abhängigkeit vom Winkel φ mit

$$\begin{aligned}x &= \cos \varphi, \\ y &= \sin \varphi\end{aligned}$$

bezeichnet.

Hierdurch wird die Sinusfunktion $\sin \varphi$ und die Cosinusfunktion $\cos \varphi$ definiert.

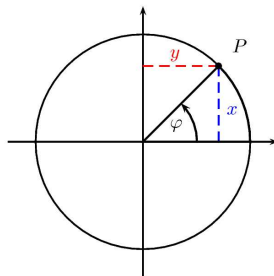
Dabei misst man den Winkel φ im Bogenmaß:

Länge des Kreissegments von $(1, 0)$ bis zum Punkt P.

Der Umfang eines Einheitskreises beträgt $2 \cdot \pi$. Dabei beschreibt π die Kreiszahl

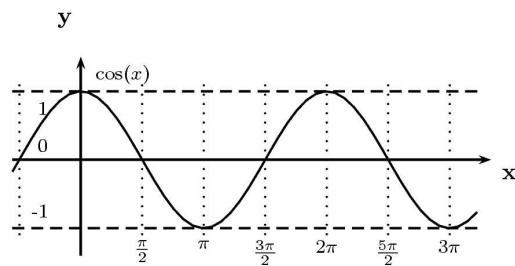
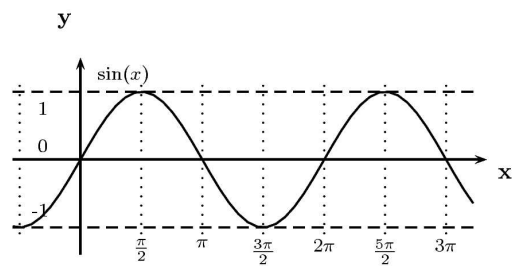
$$\pi \approx 3,14159 \dots$$

Da das Bogenmaß 2π einem Gradmaß von 360° entspricht, gilt:



Gradmaß	Bogenmaß
0°	0
45°	$\frac{\pi}{4}$
90°	$\frac{\pi}{2}$
α°	$\frac{\pi}{180^\circ} \cdot \alpha$

Sinus- und Cosinusfunktion haben folgende Gestalt:



Es gilt:

- (a) $\sin^2 \varphi + \cos^2 \varphi = 1$ (Übungsaufgabe)
- (b) $\sin(-\varphi) = -\sin \varphi$ (ungerade Fkt.)
- (c) $\cos(-\varphi) = \cos \varphi$ (gerade Fkt.)
- (d) 2π -Periodizität:

$$\begin{aligned}\sin(\varphi + 2\pi k) &= \sin \varphi & \forall k \in \mathbb{Z}, \\ \cos(\varphi + 2\pi k) &= \cos \varphi & \forall k \in \mathbb{Z}.\end{aligned}$$

(e) Additionstheoreme:

$$\begin{aligned}\cos(\alpha + \beta) &= \cos \alpha \cdot \cos \beta - \sin \alpha \cdot \sin \beta, \\ \sin(\alpha + \beta) &= \sin \alpha \cdot \cos \beta + \cos \alpha \cdot \sin \beta.\end{aligned}$$

(f) Potenzreihendarstellung:

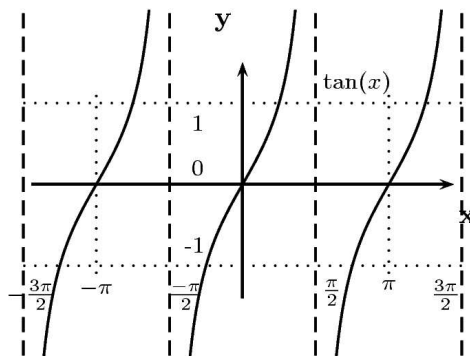
$$\begin{aligned}\sin x &= \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k+1}}{(2k+1)!} = x - \frac{x^3}{3!} + \frac{x^5}{5!} \mp \dots, \\ \cos x &= \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k}}{(2k)!} = 1 - \frac{x^2}{2!} + \frac{x^4}{4!} \mp \dots\end{aligned}$$

Diese Reihen konvergieren absolut für alle $x \in \mathbb{R}$. Die Sinus- und Cosinusfunktion sind stetig. Es gilt die Moivre-Formel:

$$e^{ix} = \cos x + i \sin x \quad \forall x \in \mathbb{R} \quad (\text{Übungsaufgabe}).$$

Die Tangensfunktion $\tan x$ ist definiert als

$$\tan \varphi = \frac{\sin \varphi}{\cos \varphi}.$$



π -Periodizität:

$$\tan(\varphi) = \tan(\varphi + k\pi) \quad \forall k \in \mathbb{Z}.$$

Sie ist stetig auf

$$\mathbb{R} \setminus \left\{ \frac{2k+1}{2}\pi \mid k \in \mathbb{Z} \right\}.$$

21.7 Trigonometrische Umkehrfunktionen

Wenn wir die trigonometrischen Funktionen auf ein Intervall einschränken, auf dem sie streng monoton sind, können wir dort Umkehrfunktionen definieren.

- (a) $\cos$ ist in $[0, \pi]$ streng monoton fallend mit Wertebereich $[-1, 1]$. Die Umkehrfunktion

$$\arccos : [-1, 1] \rightarrow [0, \pi]$$

heißt Arcus-Cosinus. Sie ist stetig.

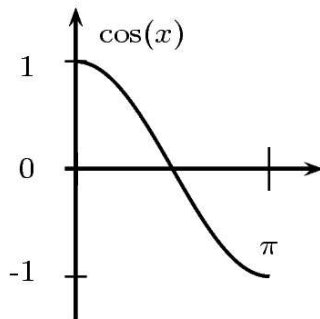


Abbildung 21.7: Cosinus

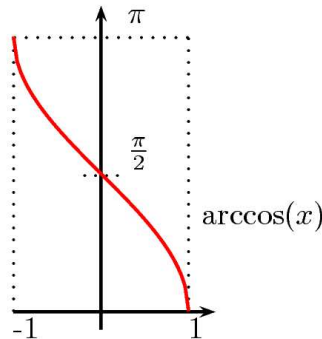


Abbildung 21.8: Arcus-Cosinus

- (b) $\sin$ ist in $[-\frac{\pi}{2}, \frac{\pi}{2}]$ streng monoton wachsend mit Wertebereich $[-1, 1]$. Die Umkehrfunktion

$$\arcsin : [-1, 1] \rightarrow [-\frac{\pi}{2}, \frac{\pi}{2}]$$

heißt Arcus-Sinus. Sie ist stetig.

- (c) $\tan$ ist in $(-\frac{\pi}{2}, \frac{\pi}{2})$ streng monoton wachsend mit Wertebereich $\mathbb{R}$. Die Umkehrfunktion

$$\arctan : \mathbb{R} \rightarrow (-\frac{\pi}{2}, \frac{\pi}{2})$$

heißt Arcus-Tangens. Sie ist stetig.

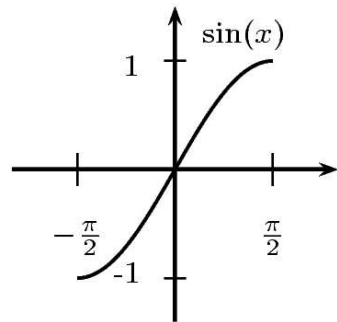


Abbildung 21.9: Sinus

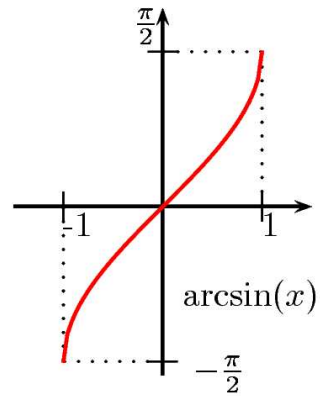


Abbildung 21.10: Arcus-Sinus

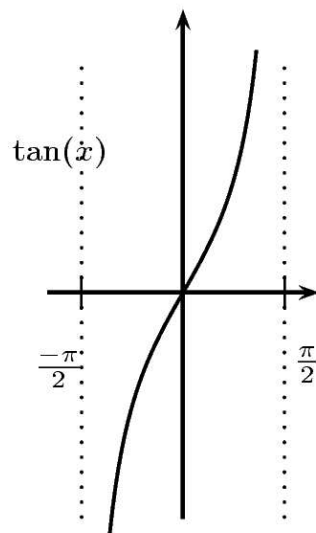


Abbildung 21.11: Tangens

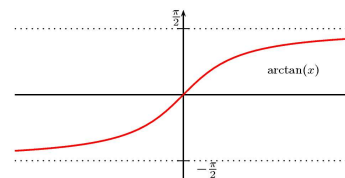


Abbildung 21.12: Arcus-Tangens

Kapitel 22

Differenzierbarkeit

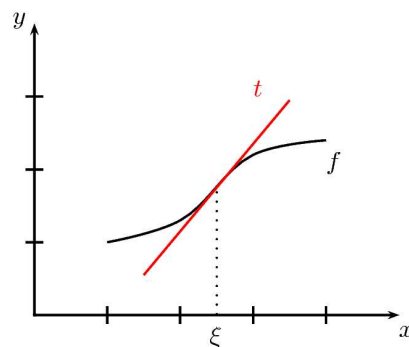
22.1 Motivation

- Wird durch eine Funktion beschrieben, wie sich eine Größe abhängig von einer anderen verändert, so stellt sich die Frage:

„Wie schnell“ ändert sich die abhängige Größe?

Beschreibt die Funktion z.B. den Ort eines Punktes abhängig v.d. Zeit (Bewegung entlang einer Linie), so ist dies die Frage nach der Geschwindigkeit (zu jedem Zeitpunkt).

- Betrachtet man eine Funktion nur in der unmittelbaren Nähe einer Stelle, dann möchte man sie dort durch eine einfachere Funktion möglichst gut annähern, z.B. $f(x)$ durch $t(x) = ax + b$ ($a, b \in \mathbb{R}$ (affin-lineare Funktion)).



Diese und ähnliche Fragen beantwortet die Ableitung (auch Differentiation genannt).

22.2 Definition: (differenzierbar, Ableitung, Differentialquotient)

Es sei eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ und ein $\xi \in \mathbb{R}$ gegeben. $f(x)$ heißt differenzierbar in ξ , falls der Grenzwert

$$\lim_{x \rightarrow \xi} \frac{f(x) - f(\xi)}{x - \xi}$$

existiert.

In diesem Fall heißt der Grenzwert Ableitung von $f(x)$ in ξ oder auch Differentialquotient von $f(x)$ in ξ und wird mit

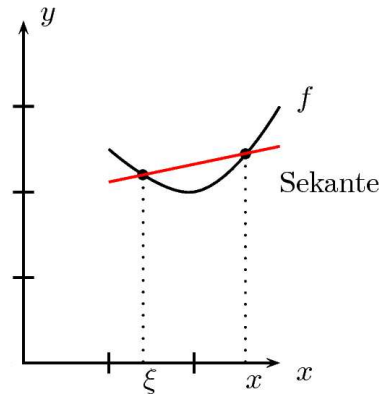
$$f'(\xi) \quad \text{oder} \quad \frac{df}{dx}(\xi)$$

bezeichnet.

Ist f in allen $\xi \in \mathbb{R}$ differenzierbar, so heißt f differenzierbar.

22.3 Bemerkungen:

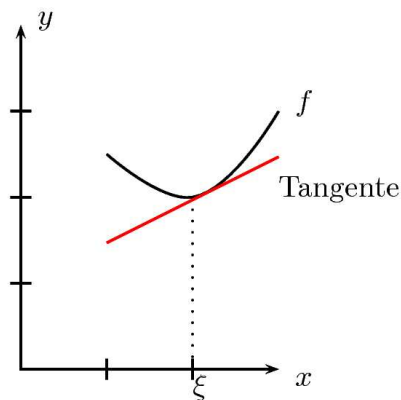
- (a) Der Ausdruck $\frac{f(x) - f(\xi)}{x - \xi} =: \frac{\Delta y}{\Delta x}$ heißt Differenzenquotient. Er gibt die Steigung der Sekante zwischen den Punkten $(\xi, f(\xi))$ und $(x, f(x))$ des Graphen von f an.



Für $x \rightarrow \xi$ (also $\Delta x \rightarrow 0$) geht die Sekantensteigung in der Tangentensteigung im Punkt $(\xi, f(\xi))$ über, $f'(\xi) = \frac{dy}{dx}(\xi) = \lim_{x \rightarrow \xi} \frac{\Delta y}{\Delta x}$.

- (b) Gleichung der Tangente t an f in $(\xi, f(\xi))$:

$$t(x) = f(\xi) + f'(\xi) \cdot (x - \xi).$$



- (c) Schränkt man bei der Grenzwertbildung x auf Werte größer/kleiner als ξ ein, so erhält man die einseitigen Grenzwerte

$$f'(\xi^+) := \lim_{x \rightarrow \xi^+} \frac{f(x) - f(\xi)}{x - \xi}$$

$$f'(\xi^-) := \lim_{x \rightarrow \xi^-} \frac{f(x) - f(\xi)}{x - \xi},$$

die als rechtsseitige bzw. linksseitige Ableitung von $f(x)$ in ξ bezeichnet werden.

- (d) Wie im Falle der Stetigkeit übertragen sich die Begriffe sinngemäß auf Funktionen mit Definitionsbereich $D \subsetneq \mathbb{R}$.

22.4 Beispiel:

Es sei $f(x) = x^n$, $n \in \mathbb{N}$, $x \in \mathbb{R}$.

Für $x, \xi \in \mathbb{R}$ gilt dann:

$$f(x) - f(\xi) = x^n - \xi^n = (x - \xi)(x^{n-1} + x^{n-2}\xi + \dots + x\xi^{n-2} + \xi^{n-1})$$

Für $x \neq \xi$ also:

$$\frac{f(x) - f(\xi)}{x - \xi} = \frac{x^n - \xi^n}{x - \xi} = x^{n-1} + x^{n-2}\xi + \dots + x\xi^{n-2} + \xi^{n-1},$$

$$\lim_{x \rightarrow \xi} \frac{f(x) - f(\xi)}{x - \xi} = n\xi^{n-1}.$$

Also ist $f(x) = x^n$ überall auf $\mathbb{R}$ differenzierbar mit $f'(x) = nx^{n-1}$.

22.5 Ableitungen elementarer Funktionen

$f(x)$	$f'(x)$
$x^\alpha, \alpha \in \mathbb{R}, x > 0$	$\alpha x^{\alpha-1}$
e^x	e^x
$\sin x$	$\cos x$
$\cos x$	$-\sin x$
$\tan x, x \neq \frac{\pi}{2} + k\pi, k \in \mathbb{Z}$	$\frac{1}{\cos^2 x}$
$\ln x, x > 0$	$\frac{1}{x}$

Um aus diesen Ableitungen viele weitere Funktionen „zusammensetzen“ zu können, benötigt man die Ableitungsregeln, die im folgenden Satz zusammengefasst sind.

22.6 Satz: (Differentiationsregeln)

(a) Ist $f : \mathbb{R} \rightarrow \mathbb{R}$ in $\xi \in \mathbb{R}$ differenzierbar, so ist f in ξ auch stetig.

Sind $f : \mathbb{R} \rightarrow \mathbb{R}$ und $g : \mathbb{R} \rightarrow \mathbb{R}$ in ξ differenzierbar, so gilt:

(b) $\alpha f + \beta g$ ist differenzierbar für $\alpha, \beta \in \mathbb{R}$ mit

$$(\alpha f + \beta g)'(\xi) = \alpha f'(\xi) + \beta g'(\xi).$$

(c) $f \cdot g$ ist differenzierbar in ξ mit

$$(f \cdot g)'(\xi) = f'(\xi) \cdot g(\xi) + f(\xi) \cdot g'(\xi). \quad (\text{Produktregel})$$

(d) Falls $g(\xi) \neq 0$, so ist $\frac{f}{g}$ in ξ differenzierbar mit

$$\left(\frac{f}{g}\right)'(\xi) = \frac{f'(\xi) \cdot g(\xi) - f(\xi) \cdot g'(\xi)}{(g(\xi))^2}. \quad (\text{Quotientenregel})$$

(e) Ist $f : \mathbb{R} \rightarrow \mathbb{R}$ in $\xi \in \mathbb{R}$ und $g : \mathbb{R} \rightarrow \mathbb{R}$ in $\eta = f(\xi) \in \mathbb{R}$ differenzierbar, so ist auch die Komposition $(g \circ f)(x) = g(f(x))$ in ξ differenzierbar mit

$$(g \circ f)'(\xi) = g'(f(\xi)) \cdot f'(\xi). \quad (\text{Kettenregel})$$

(f) Ist $f : [a, b] \rightarrow \mathbb{R}$ streng monoton wachsend und in $\xi \in [a, b]$ differenzierbar mit $f'(\xi) \neq 0$, so ist auch die inverse Funktion

$$f^{-1} : [f(a), f(b)] \rightarrow \mathbb{R}$$

in $\eta = f(\xi)$ differenzierbar mit

$$(f^{-1})'(\nu) = \frac{1}{f'(\xi)}.$$

Bemerkung: Auch diese Aussagen übertragen sich sinngemäß auf Funktionen mit Definitionsbereich $D \subseteq \mathbb{R}$. Aussage (f) gilt auch für streng monoton fallende Funktionen (Definitionsbereich von f^{-1} dann $[f(b), f(a)]$).

Beweis:

Wir zeigen nur beispielhaft (c) (Produktregel). Die übrigen Beweise verlaufen ähnlich.

Aufgrund der Grenzwertsätze 20.4 auf Seite 150 gilt

$$\begin{aligned} \lim_{x \rightarrow \xi} \frac{f(x)g(x) - f(\xi)g(\xi)}{x - \xi} &= \lim_{x \rightarrow \xi} \frac{f(x)g(x) - \overbrace{f(\xi)g(x) + f(\xi)g(x) - f(\xi)g(\xi)}^{\text{„geschickt Null addiert“}}}{x - \xi} \\ &= \lim_{x \rightarrow \xi} \left(\frac{f(x) - f(\xi)}{x - \xi} g(x) + f(\xi) \frac{g(x) - g(\xi)}{x - \xi} \right) \\ &= f'(\xi)g(\xi) + f(\xi)g'(\xi). \end{aligned}$$

□

22.7 Beispiele

(a) Mit $\frac{d}{dx}(\sin x) = \cos x$ und $\frac{d}{dx}(\cos x) = -\sin x$ folgt:

$$\begin{aligned} \frac{d}{dx}(\tan x) &= \frac{d}{dx} \left(\frac{\sin x}{\cos x} \right) \\ &\stackrel{20.6 \text{ auf Seite 150. (c)}}{=} \frac{\cos x \cos x - \sin x(-\sin x)}{\cos^2 x} \\ &= \frac{\cos^2 x + \sin^2 x}{\cos^2 x} \\ &= \frac{1}{\cos^2 x} \end{aligned}$$

für $x \neq (2k+1)\frac{\pi}{2}$, $k \in \mathbb{Z}$.

(b) Mit $y = \tan x$ erhält man mit 22.6 auf der vorherigen Seite.(f) als Ableitung des Arcustangens

$$\begin{aligned} \frac{d}{dy}(\arctan y) &= \frac{1}{\frac{d}{dx}(\tan x)} \\ &= \frac{1}{\frac{1}{\cos^2 x}} = \cos^2 x \\ &= \frac{\cos^2 x}{\sin^2 x + \cos^2 x} = \frac{1}{1 + \frac{\sin^2 x}{\cos^2 x}} = \frac{1}{1 + \tan^2} \\ &= \frac{1}{1 + y^2}. \end{aligned}$$

(c) $\frac{d}{dx}(\sin(e^x)) = \cos(e^x) \cdot e^x.$

(d) $\frac{d}{dx}(a^x) = \frac{d}{dx}(e^{(\ln a) \cdot x}) = e^{(\ln a) \cdot x} \ln a = a^x \ln a \quad (a > 0).$

(e) Sei $x > 0$.

$$\begin{aligned} \frac{d}{dx}(x^x) &= \frac{d}{dx}(e^{(\ln x)x}) = e^{(\ln x)x} \left((\ln x) \cdot 1 + \frac{1}{x} \cdot x \right) \\ &= x^x (\ln x + 1). \end{aligned}$$

22.8 Bemerkung:

Während Differenzierbarkeit Stetigkeit impliziert, gilt die Umkehrung nicht:

$f(x) = |x|$ ist in $x = 0$ stetig, dort aber nicht differenzierbar
($f'(0^+) = 1 \neq -1 = f'(0^-)$).

22.9 Definition: (Zweite Ableitung)

Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ differenzierbar mit Ableitung f' . Falls $f'(x)$ wiederum differenzierbar ist, so erhält man die zweite Ableitung

$$f''(x) := f^{(2)}(x) := \frac{d^2}{dx^2} f(x).$$

Analog definiert man höhere Ableitungen $f'''(x), f^{(4)}(x), \dots, f^{(n)}(x)$. Ist $f(x)$ n -mal differenzierbar auf $\mathbb{R}$ und die n -te Ableitung $\frac{d^n}{dx^n} f(x)$ stetig auf $\mathbb{R}$, so schreibt man $f \in C^n(\mathbb{R})$ und bezeichnet f als n -mal stetig differenzierbar auf $\mathbb{R}$. Gilt dies für alle $n \in \mathbb{N}$, schreibt man $f \in C^\infty(\mathbb{R})$.

Bemerkung: Die Definitionen übertragen sich sinngemäß auf Intervalle $[a, b]$, wobei in a nur rechtseitige und in b nur linksseitige Ableitungen zugelassen sind. Ist f auf $[a, b]$ n -mal stetig differenzierbar, schreibt man $f \in C^n[a, b]$. Statt $C^0[a, b]$ oder $C^0(\mathbb{R})$ schreibt man $C[a, b]$ bzw. $C(\mathbb{R})$ für die stetigen Funktionen auf $[a, b]$ bzw. $\mathbb{R}$.

22.10 **Beispiel:**

$$\begin{aligned}f(x) &= 5x^3 + 6x^2 - 2, \\f'(x) &= 15x^2 + 12x, \\f''(x) &= 30x + 12, \\f^{(3)}(x) &= 30, \\f^{(n)}(x) &= 0 \quad \forall n \in \mathbb{N} \text{ mit } n \geq 4, \\f &\in C^\infty(\mathbb{R}).\end{aligned}$$

Kapitel 23

Mittelwertsätze und die Regel von L'Hospital

23.1 Motivation

- Für stetige Funktionen gibt es wichtige Aussagen wie den Nullstellensatz 20.9 (a) und den Zwischenwertsatz 20.9 (b). Gibt es ähnliche Aussagen für differenzierbare Funktionen?
- Gibt es einen Trick, wie man Grenzwerte der Form „ $\frac{0}{0}$ “ oder „ $\frac{\infty}{\infty}$ “ einfach berechnen kann?

23.2 Satz: (Mittelwertsätze)

(a) Satz von Rolle:

Sei $f \in C^1[a, b]$ mit $f(a) = f(b)$. Dann existiert ein $\xi \in (a, b)$ mit $f'(\xi) = 0$.

(b) (Erster) Mittelwertsatz:

Sei $f \in C^1[a, b]$. Dann existiert ein $\xi \in (a, b)$ mit

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

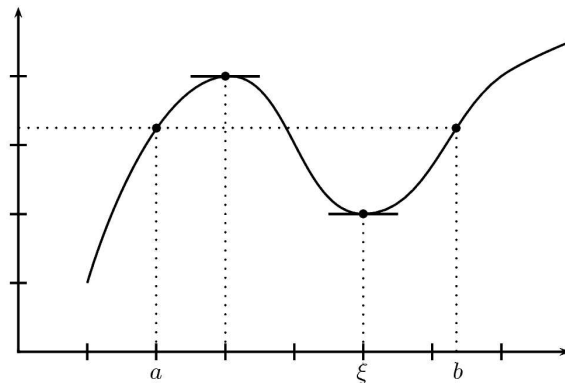
(c) Zweiter Mittelwertsatz:

Seien $f, g \in C^1[a, b]$ mit $g'(x) \neq 0 \forall x \in (a, b)$. Dann existiert ein $\xi \in (a, b)$ mit

$$\frac{f'(\xi)}{g'(\xi)} = \frac{f(b) - f(a)}{g(b) - g(a)}.$$

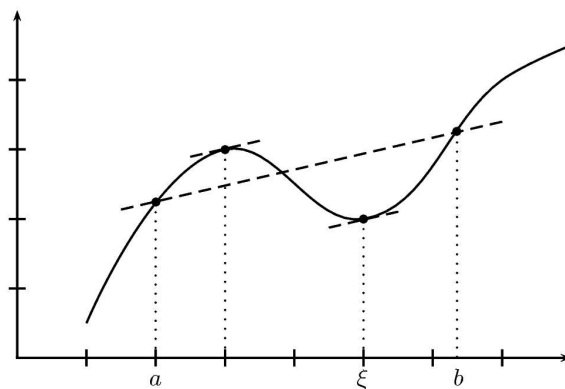
Veranschaulichung:

(a) Satz von Rolle:



Es gibt eine waagerechte Tangente in (a, b) , falls $f(a) = f(b)$ (Hier gibt es sogar zwei Stellen ξ mit $f'(\xi) = 0$).

(b) Mittelwertsatz:



Zu jeder Sekante existiert mindestens eine parallele Tangente (hier: 2).

Beweis:

Wir zeigen nur (b).

Für $f \in C^1[a, b]$ erfüllt

$$h(x) := f(x) - \frac{x-a}{b-a} (f(b) - f(a))$$

die Voraussetzung des Satzes von Rolle:

$$\left. \begin{array}{l} h \in C^1[a, b] \\ h(a) = f(a) \\ h(b) = f(b) - \frac{b-a}{b-a} (f(b) - f(a)) = f(a) \end{array} \right\} \Rightarrow h(a) = h(b).$$

$$\Rightarrow \exists \xi \in (a, b) \text{ mit } 0 = h'(\xi) = f'(\xi) - \frac{1}{b-a} (f(b) - f(a)).$$

$$\Rightarrow f'(\xi) = \frac{f(b) - f(a)}{b - a}.$$

□

Der zweite Mittelwertsatz hat eine interessante Anwendung:

23.3 Satz: (Regel von l'Hospital)

(a) Fall „ $\frac{0}{0}$ “:

Seien f, g stetig differenzierbar auf (a, b) , $\xi \in (a, b)$, $f(\xi) = g(\xi) = 0$ und es gelte $g'(x) \neq 0$ für $x \neq \xi$. Dann folgt

$$\lim_{x \rightarrow \xi} \frac{f(x)}{g(x)} = \lim_{x \rightarrow \xi} \frac{f'(x)}{g'(x)}$$

wenn der rechte Grenzwert existiert.

(b) Fall „ $\frac{\infty}{\infty}$ “:

Seien f, g stetig differenzierbar auf $(a, b) \setminus \xi$, und

$$\lim_{x \rightarrow \xi} f(x) = \lim_{x \rightarrow \xi} g(x) = \infty$$

und es gelte $g'(x) \neq 0$ für $x \neq \xi$. Dann folgt

$$\lim_{x \rightarrow \xi} \frac{f(x)}{g(x)} = \lim_{x \rightarrow \xi} \frac{f'(x)}{g'(x)}$$

wenn der rechte Grenzwert existiert.

Beweis:

von (a)

Wegen $f(\xi) = g(\xi) = 0$ und dem 2. Mittelwertsatz existiert $\eta = \eta(x)$ mit

$$\frac{f(x)}{g(x)} = \frac{f(x) - f(\xi)}{g(x) - g(\xi)} \stackrel{2.MWS}{=} \frac{f'(\eta)}{g'(\eta)}$$

mit $\eta(x)$ zwischen x und ξ . Für $x \rightarrow \xi$ gilt daher auch $\eta(x) \rightarrow \xi$ und es folgt die Behauptung.

□

23.4 Beispiel:

- (a) $\lim_{x \rightarrow \infty} \frac{\sin x}{x}$ hat die Form „ $\frac{0}{0}$ “. Mit l'Hospital ergibt sich

$$\lim_{x \rightarrow \infty} \frac{\sin x}{x} = \lim_{x \rightarrow \infty} \frac{\cos x}{1} = \cos 0 = 1.$$

- (b) $\lim_{x \rightarrow \infty} \frac{x}{e^x}$ ist vom Typ „ $\frac{\infty}{\infty}$ “. l'Hospital liefert

$$\lim_{x \rightarrow \infty} \frac{x}{e^x} = \lim_{x \rightarrow \infty} \frac{1}{e^x} = \lim_{x \rightarrow \infty} e^{-x} = 0.$$

Kapitel 24

Der Satz von Taylor

24.1 Motivation

- Für eine differenzierbare Funktion $f(x)$ stellt die Tangente

$$t(x) = f(\xi) + (x - \xi) \cdot f'(\xi)$$

eine lokale Approximation durch ein Polynom 1. Grades im Punkt ξ dar.

- Ist es möglich, $f(x)$ in ξ durch ein Polynom höheren Grades zu approximieren, wenn f eine höhere Differenzierbarkeitsordnung besitzt?

24.2 Satz: (Satz von Taylor)

Sei $\xi \in (a, b)$ und $f \in C^{m+1}[a, b]$. Dann besitzt $f(x)$ folgende Taylorentwicklung um ξ :

$$f(x) = T_m(x, \xi) + R_m(x, \xi)$$

mit dem Taylorpolynom m -ten Grades

$$T_m(x, \xi) = \sum_{k=0}^m \frac{(x - \xi)^k}{k!} f^{(k)}(\xi)$$

und dem Restglied nach Lagrange

$$R_m(x, \xi) = \frac{(x - \xi)^{m+1}}{(m+1)!} f^{(m+1)}(\xi + \Theta(x - \xi))$$

mit $\Theta \in (0, 1)$.

Beweis:

Betrachte $g(x) := f(x) - T_m(x, \xi)$.

Dann ist

$$\begin{aligned} g(\xi) &= f(\xi) - T_m(\xi, \xi) \\ &= f(\xi) - f(\xi) = 0, \\ g'(\xi) &= f'(\xi) - T'_m(\xi, \xi) \\ &= f'(\xi) - f'(\xi) = 0, \\ &\vdots \\ g^{(m)}(\xi) &= f^{(m)}(\xi) - f^{(m)}(\xi) = 0. \end{aligned}$$

Nun wenden wir den 2. Mittelwertsatz auf $\frac{g(x)}{(x - \xi)^{m+1}}$ an:
 $\exists \xi_1$ zwischen x und ξ mit

$$\begin{aligned} \frac{g(x)}{(x - \xi)^{m+1}} &= \frac{g(x) - \overbrace{g(\xi)}^0}{(x - \xi)^{m+1} - (\xi - \xi)^{m+1}} \\ &\stackrel{\text{2.MWS}}{=} \frac{g'(\xi)}{(m+1)(\xi_1 - \xi)^m}. \end{aligned}$$

Induktiv folgt

$$\begin{aligned} \frac{g(x)}{(x - \xi)^{m+1}} &= \frac{g'(\xi_1)}{(m+1)(\xi_1 - \xi)^m} \\ &= \frac{g''(\xi_2)}{m \cdot (m+1) \cdot (\xi_2 - \xi)^{m-1}} \\ &= \dots = \frac{g^{(m)}(\xi_m)}{(m+1)!(\xi_m - \xi)} \\ &= \frac{g^{(m+1)}(\xi_{m+1})}{(m+1)!} \end{aligned}$$

mit ξ_k zwischen x und ξ für $k = 1, \dots, m+1$. Mit $g(x) = f(x) - T_m(x, \xi)$ folgt also:

$$f(x) - T_m(x, \xi) = \frac{(x - \xi)^{m+1}}{(m+1)!} g^{(m+1)}(\xi_{m+1}).$$

Wegen $g^{(m+1)}(x) = f^{(m+1)}(x)$ ist daher

$$f(x) = T_m(x, \xi) + \frac{(x - \xi)^{m+1}}{(m+1)!} f^{(m+1)}(\xi_{m+1})$$

mit ξ_{m+1} zwischen x und ξ .

□

24.3 Bemerkung:

- (a) Für $m = 0$ liefert der Satz von Taylor den 1. Mittelwertsatz.
- (b) Man kann sogar zeigen, dass $T_m(x, \xi)$ das einzigste Polynom vom Grad $\leq m$ ist, das die Approximationsgüte $O\left((x - \xi)^{m+1}\right) = o\left((x - \xi)^m\right)$ besitzt.
- (c) Neben der Restglieddarstellung nach Lagrange gibt es noch weitere Darstellungen des Restglieds.

24.4 Beispiel:

- (a) Taylor-Entwicklung der Exponentialfunktion um $\xi = 0$:

Aus $f(x) = \sum_{k=0}^m \frac{(x - \xi)^k}{k!} f^{(k)}(\xi) + \frac{(x - \xi)^{m+1}}{(m+1)!} f^{(m+1)}(\xi + \Theta(x - \xi))$ folgt
mit $\frac{d^k}{dx^k}(e^x) = e^x$; $e^0 = 1$ und $\xi = 0$:

$$f(x) = \sum_{k=0}^m \frac{x^k}{k!} + \underbrace{\frac{x^{m+1}}{(m+1)!} e^{\Theta x}}_{R_m(x,0)} \quad \text{für } 0 < \Theta < 1.$$

Für $0 \leq x \leq 1$ hat man beispielsweise die Fehlerabschätzung

$$|R_m(x, 0)| \leq \frac{x^{m+1}}{(m+1)!} e^{\Theta x} \leq \frac{1}{(m+1)!} \cdot e.$$

Hieraus folgt für $m = 10$:

$$|R_{10}(x, 0)| \leq \frac{1}{11!} \cdot e \approx 6.81 \cdot 10^{-8}.$$

- (b) Taylor-Entwicklung der Sinusfunktion um $\xi = 0$:

Mit $\frac{d}{dx} \sin x = \cos x$, $\frac{d}{dx} \cos x = -\sin x$ und $\sin 0 = 0$, $\cos 0 = 1$ folgt, dass $T_m(x, 0)$ keine geraden Potenzen in x enthält:

$$\begin{aligned} \frac{d}{dx}(\sin x) &= \cos x & \cos 0 &= 1 \\ \frac{d^2}{dx^2}(\sin x) &= -\sin x & -\sin 0 &= 0 \\ \frac{d^3}{dx^3}(\sin x) &= -\cos x & -\cos 0 &= -1 \\ \frac{d^4}{dx^4}(\sin x) &= \sin x & \sin 0 &= 0 \\ \frac{d^5}{dx^5}(\sin x) &= \cos x & \cos 0 &= 1 \\ &\vdots & & \end{aligned}$$

Damit erhalten wir folgende Taylorentwicklung für $m = 2n + 1$:

$$\begin{aligned}\sin x &= \sum_{k=0}^{2n+1} \frac{x^k}{k!} \frac{d^k}{dx^k}(\sin x) \Big|_{x=0} + R_{2n+2}(x, 0) \\ &= x - \frac{x^3}{3!} + \frac{x^5}{5!} \mp \dots + (-1)^n \frac{x^{2n+1}}{(2n+1)!} + R_{2n+2}(x, 0)\end{aligned}$$

mit $R_{2n+2}(x, 0) = (-1)^{n+1} \frac{\cos(\Theta x)}{(2n+3)!} x^{2n+3}$ $0 < \Theta < 1$.

Für $|x| \leq 1$ und $n = 3$ folgt beispielsweise

$$|R_8(x, 0)| \leq \frac{1}{9!} \approx 2,8 \cdot 10^{-6}.$$

24.5 Definition: (Taylorreihen)

Für eine C^∞ -Funktion f bezeichnet man die Potenzreihe $\sum_{k=0}^{\infty} \frac{(x-\xi)^k}{k!} f^{(k)}(\xi)$ als die Taylorreihe von f im Entwicklungspunkt ξ .

24.6 Bemerkungen:

- (a) Die Taylorreihe muss i.A. nicht konvergieren!
- (b) Konvergiert sie, so muss sie nicht gegen $f(x)$ konvergieren.
- (c) Ist dies jedoch der Fall, so dass also

$$f(x) = \sum_{k=0}^{\infty} \frac{(x-\xi)^k}{k!} f^{(k)}(\xi) \quad \forall x \in (a, b) \text{ gilt,}$$

so heißt f reell analytisch (C^ω -Funktion) auf (a, b) .

- (d) Die Taylorreihe einer C^∞ -Funktion f mit Entwicklungspunkt ξ konvergiert in x genau dann gegen $f(x)$, wenn

$$\lim_{m \rightarrow \infty} R_m(x, \xi) = 0.$$

24.7 Beispiele

- (a) $\sum_{k=0}^{\infty} \frac{x^k}{k!}$ ist die Taylorreihe zur Exponentialfunktion im Entwicklungspunkt $\xi = 0$. Sie konvergiert für alle $x \in \mathbb{R}$ gegen e^x .

Allgemein gilt: Wenn eine Funktion eine Potenzreihendarstellung hat, so ist dies die Taylorreihe.

(b) Die Binomialreihe $\sum_{k=0}^{\infty} \binom{\alpha}{k} x^k$ ist die Taylorreihe von $(1+x)^\alpha$ im Entwicklungspunkt $\xi = 0$. Sie konvergiert für $|x| < 1$ (vgl. 19.11 auf Seite 147).

(c) $f(x) = \begin{cases} e^{-\frac{1}{x}} & \text{für } x > 0 \\ 0 & \text{sonst} \end{cases}$ ist aus $C^\infty(\mathbb{R})$.

Nachrechnen zeigt, dass

$$f^{(m)}(0) = 0 \quad \forall m \in \mathbb{N}_0.$$

$$\Rightarrow T_m(x, 0) = 0 \quad \forall m \in \mathbb{N}.$$

Die Taylorreihe verschwindet also ($f(x) = R_m(x, 0)$).

Für $x > 0$ konvergiert die Taylorreihe nicht gegen $f(x)$. $f(x)$ ist nicht reell analytisch.

24.8 Satz: (Gliederweise Differentiation von Potenzreihen)

Die durch die Potenzreihe $\sum_{k=0}^{\infty} a_k x^k$ innerhalb ihres Konvergenzradius r dargestellte Funktion ist gliederweise differenzierbar.

Die Potenzreihe

$$\sum_{k=0}^{\infty} k \cdot a_k \cdot x^{k-1}$$

hat den selben Konvergenzradius und stellt in $(-r, r)$ die Funktion dar.

24.9 Beispiel:

Die Sinusreihe $\sum_{k=0}^{\infty} (-1)^k \frac{x^{2k+1}}{(2k+1)!}$ konvergiert für alle $x \in \mathbb{R}$ gegen $\sin x$. Gliederweise Differentiation ergibt

$$\sum_{k=0}^{\infty} (-1)^k \frac{2k+1}{(2k+1)!} x^{2k} = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k}}{(2k)!}.$$

Dies ist gerade die Cosinusreihe (vgl. 21.6 auf Seite 160). Sie konvergiert ebenfalls für alle $x \in \mathbb{R}$.

Kapitel 25

Geometrische Bedeutung der Ableitungen

25.1 Motivation

- Kann man aus der Kenntnis der Ableitung Aussagen über den Graphen der Funktion gewinnen?
- Die Ableitungen liefern sogar viele Aussagen:
Monotonie, Krümmungsverhalten, Extrema, Wendepunkte.
- Diese Aussagen bilden die Grundlagen der Kurvendiskussion und sind wichtig bei Optimierungsproblemen.

25.2 Definition: (Monotonie)

Sei $I \subset \mathbb{R}$ ein Intervall.

Eine Funktion $f : I \rightarrow \mathbb{R}$ heißt monoton wachsend, wenn für alle $x_1, x_2 \in I$ mit $x_2 > x_1$ gilt: $f(x_2) \geq f(x_1)$.

f ist streng monoton wachsend, wenn stets $f(x_2) > f(x_1)$ gilt.

Die Eigenschaften monoton fallend bzw. streng monoton fallend sind durch $f(x_2) \leq f(x_1)$ bzw. $f(x_2) < f(x_1)$ definiert.

25.3 Satz: (Monotonie differenzierbarer Funktionen)

Sei f differenzierbar im Intervall I . Dann gilt:

- (a) f ist monoton wachsend (monoton fallend) in I
 $\Leftrightarrow f'(x) \geq 0$ ($f'(x) \leq 0$) in I .
- (b) $f'(x) > 0$ ($f'(x) < 0$)
 $\Rightarrow f$ ist streng monoton wachsend (streng monoton fallend) in I .

Beweis:

- (a) „ $\Rightarrow$ “: Sei f monoton wachsend und differenzierbar in I .

$$\Rightarrow f(x+h) - f(x) \geq 0 \text{ für } h > 0.$$

$$f(x+h) - f(x) \leq 0 \text{ für } h < 0.$$

für alle $x, x+h \in I$.

$$\Rightarrow \frac{f(x+h) - f(x)}{h} \geq 0 \quad \forall h \neq 0 \text{ mit } x+h \in I.$$

$$\Rightarrow f'(x) \geq 0.$$

Analog zeigt man $f'(x) \leq 0$ für f monoton fallend.

„ $\Leftarrow$ “: Sei $f'(x) \geq 0 \quad \forall x \in I$. Sei $h > 0$ und $x_0, x_0+h \in I$.

Nach dem Mittelwertsatz existiert $\Theta \in (0, 1)$ mit

$$\frac{f(x_0+h) - f(x_0)}{h} = f'(x_0 + \Theta h) \geq 0.$$

$$\Rightarrow f(x_0+h) \geq f(x_0) \text{ für } x_0+h > x_0.$$

d.h. f ist monoton wachsend.

Der Fall $f'(x) \leq 0$ wird analog behandelt.

- (b) wie (a) „ $\Leftarrow$ “, wobei $\geq$ durch $>$ ersetzt wird.

□

25.4 Beispiele

- (a) Monotonie von $f(x) = \frac{x}{x^2 + 2x + 3}$?

$$f'(x) = \frac{(x^2 + 2x + 3) \cdot 1 - (2x + 2) \cdot x}{\underbrace{(x^2 + 2x + 3)^2}_{\text{Nenner} > 0}} = \frac{-x^2 + 3}{(x^2 + 2x + 3)^2}$$

$$f'(x) \geq 0 \text{ für } 3 - x^2 \geq 0 \quad \text{d.h.} \quad |x| \leq \sqrt{3} \quad \text{d.h.} \quad x \in [-\sqrt{3}, \sqrt{3}]$$

$$f'(x) \leq 0 \text{ für } 3 - x^2 \leq 0 \quad \text{d.h.} \quad x \leq -\sqrt{3} \text{ oder } x \geq \sqrt{3}.$$

Also ist $f(x)$ monoton wachsend in $[-\sqrt{3}, \sqrt{3}]$ und monoton fallend für $x \leq -\sqrt{3}$ oder $x \geq \sqrt{3}$.

Für $x \in (-\sqrt{3}, \sqrt{3})$ ist $f'(x) > 0$, d.h. $f(x)$ streng monoton wachsend in $(-\sqrt{3}, \sqrt{3})$.

- (b) Die Umkehrung von Satz 25.3 (b) gilt nicht!

Bsp.: $f(x) = x^3$ ist streng monoton wachsend, aber $f'(0) = 0$.

Die folgenden Begriffe sind wichtig bei Optimierungsproblemen.

25.5 Definition: (Konvexität, Konkavität)

Eine Funktion $f : I \rightarrow \mathbb{R}$ heißt konvex in I , wenn für alle $a, b \in I$ und $\Theta \in (0, 1)$ gilt:

$$f(\Theta a + (1 - \Theta)b) \leq \Theta f(a) + (1 - \Theta)f(b).$$

Gilt stattdessen

$$f(\Theta a + (1 - \Theta)b) \geq \Theta f(a) + (1 - \Theta)f(b),$$

so heißt f konkav. Gilt $<$ statt $\leq$ (bzw. $>$ statt $\geq$), so heißt f streng konvex (bzw. streng konkav).

25.6 Veranschaulichung:

Konvexe Funktionen verlaufen unterhalb ihrer Sekanten und weisen eine Linkskrümmung auf.

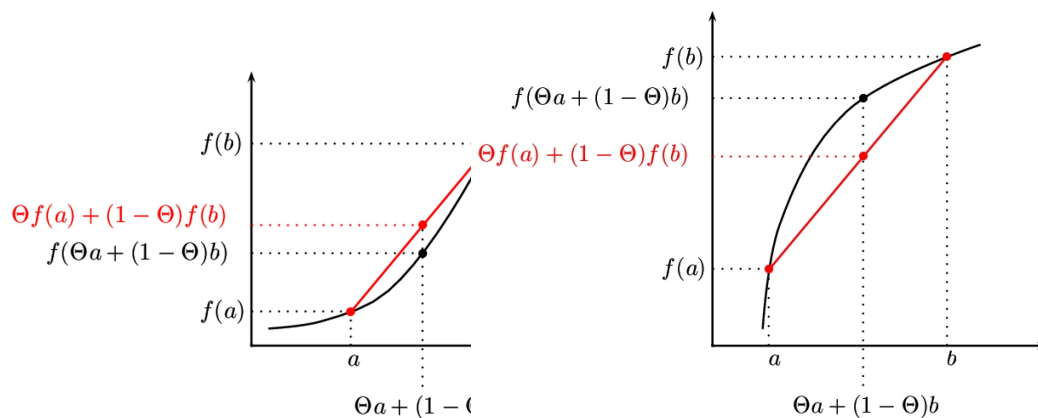


Abbildung 25.1: konvex (Linkskrümmung)

Abbildung 25.2: konkav (Rechtskrümmung)

Konkave Funktionen verlaufen oberhalb ihrer Sekanten und weisen eine Rechtskrümmung auf.

Dabei ist anschaulich klar:

25.7 Satz: (Konvexität / Konkavität von C^1 -Funktionen)

Sei $f \in C^1(I)$. Dann gilt:

- (a) f ist in I konvex $\Leftrightarrow f'$ in I monoton wachsend.

(b) f ist in I konkav $\Leftrightarrow f'$ in I monoton fallend.

Als Folgerung aus Satz 25.3 auf Seite 183 und Satz 25.7 auf der vorherigen Seite ergibt sich:

25.8 Satz: (Konverxität / Konkavität von C^2 -Funktionen)

Sei $f \in C^2(I)$. Dann gilt:

(a) f ist in I konvex $\Leftrightarrow f''(x) \geq 0$ in I .
 $f''(x) > 0$ in $I \Rightarrow f$ ist in I streng konvex.

(b) f ist in I konkav $\Leftrightarrow f''(x) \leq 0$ in I .
 $f''(x) < 0$ in $I \Rightarrow f$ ist in I streng konkav.

25.9 Beispiele

- (a) Typische konvexe Funktionen:
 $f(x) = x^2$ auf $\mathbb{R}$: $f'(x) = 2x$ monoton wachsend.
 $g(x) = e^x$ auf $\mathbb{R}$: $g'(x) = e^x$ monoton wachsend.
Wegen $f''(x) = 2 > 0$ und $g''(x) = e^x > 0$ sind f und g sogar streng konvex.
- (b) $f(x) = \ln(x)$ ist streng konkav für $x > 0$:
 $f'(x) = \frac{1}{x}$, $f''(x) = -\frac{1}{x^2} < 0$ für alle $x > 0$.

Lokale Extrema spielen eine wichtige Rolle bei Funktionen.

25.10 Definition: (Lokale Minima und Maxima)

Sei $f : I \rightarrow \mathbb{R}$ eine Funktion und ξ ein innerer Punkt in I (kein Randpunkt!). ξ heißt lokales Maximum von f in I , falls ein $\epsilon > 0$ existiert mit

$$|x - \xi| < \epsilon \Rightarrow f(x) \leq f(\xi).$$

Gilt hingegen

$$|x - \xi| < \epsilon \Rightarrow f(x) \geq f(\xi),$$

so heißt ξ lokales Minimum von f in I .

Gilt Gleichheit nur im Punkt ξ , so liegt ein strenges lokales Maximum bzw. strenges lokales Minimum vor.

Wie kann man solche lokalen Extrema finden?

25.11 Satz: (Notwendige Bedingung für lokale Extrema)

Sei f in ξ differenzierbar und habe dort ein lokales Extremum (d.h. lokales Maximum oder Minimum). Dann ist $f'(\xi) = 0$.

Beweis:

Wir nehmen an, dass in ξ ein lokales Maximum vorliegt.

$$\Rightarrow \exists \epsilon > 0 : |x - \xi| < \epsilon \Rightarrow f(x) \leq f(\xi).$$

$$\text{Falls } x > \xi: \overbrace{\frac{f(x) - f(\xi)}{x - \xi}}^{\leq 0} \leq 0 \quad (*)$$

$$\text{Falls } x < \xi: \overbrace{\frac{f(x) - f(\xi)}{x - \xi}}^{\substack{>0 \\ \leq 0}} \geq 0 \quad (**).$$

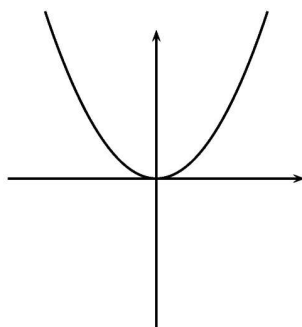
Da f in ξ differenzierbar ist, existiert ein eindeutiger Grenzwert für $x \rightarrow \xi$.

Wegen (*) und (**) gilt $f'(\xi) = 0$.

Der Beweis für ein lokales Minimum verläuft analog.

□

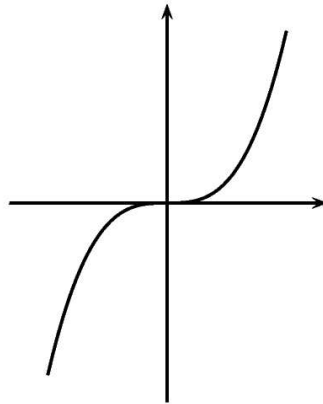
25.12 Beispiele



- (a) $f(x) = x^2$ hat in $\xi = 0$ ein (lok.) Minimum. Wegen $f'(x) = 2x$ gilt dort $f'(0) = 0$.

- (b) Die Umkehrung von Satz 25.11 gilt nicht:

Für $f(x) = x^3$ gilt $f'(x) = 3x^2$ und somit $f'(0) = 0$. Dennoch liegt in $\xi = 0$ kein lokales Extremum vor.



Satz 25.11 liefert also nur eine notwendige Bedingung, die jedoch nicht hinreichend für das Vorliegen eines lokalen Extremums ist.

25.13 Satz: (Hinreichende Bedingung für lokale Extrema)

Sei $f \in C^n(I)$, $n \geq 2$,
 ξ innerer Punkt von I , und es gelte $f'(\xi) = f''(\xi) = \dots = f^{(n-1)}(\xi) = 0$ und
 $f^{(n)}(\xi) \neq 0$. Dann gilt:

- (a) Ist n gerade, so liegt ein lokales Extremum vor.
 Ist $f^{(n)}(\xi) > 0$, ist es ein strenges lokales Minimum.
 Ist $f^{(n)}(\xi) < 0$, ist es ein strenges lokales Maximum.
- (b) Für n ungerade liegt kein Extremum vor.

Beweis:

Wir betrachten nur den Fall $f^{(n)}(\xi) > 0$. Der Fall $f^{(n)}(\xi) < 0$ geht analog. Sei h so klein, dass $\xi + h$ innerer Punkt von I ist. Nach dem Satz von Taylor existiert ein $\Theta \in (0, 1)$ mit

$$\begin{aligned} f(\xi + h) &= \sum_{k=0}^{n-2} \frac{h^k}{k!} f^{(k)}(\xi) + \frac{h^{n-1}}{(n-1)!} f^{(n-1)}(\xi + (1-\Theta)h) \\ &= f(\xi) + \sum_{k=1}^{n-1} \frac{h^k}{k-1} \cdot \underbrace{f^{(k)}(\xi)}_0 + \frac{h^{n-1}}{(n-1)!} f^{(n-1)}(\xi + (1-\Theta)h) \quad (*) \end{aligned}$$

Da $f^{(n-1)}$ stetig in I ist und $f^{(n)}(\xi) > 0$, ist $f^{(n-1)}$ in einer Umgebung um ξ streng monoton wachsend. h soll so klein sein, dass $h + \xi$ in dieser Umgebung liegt. Dann gilt:

$$f^{(n-1)}(\xi + (1-\Theta)h) \begin{cases} > 0 & (h > 0) \\ < 0 & (h < 0) \end{cases}$$

Ist n ungerade, so ist $h^{n-1} > 0$ und somit

$$h^{n-1} f^{(n-1)}(\xi + (1 - \Theta)h) \begin{cases} > 0 & (h > 0) \\ < 0 & (h < 0) \end{cases}.$$

Wegen (*) kann also kein strenges Extremum vorliegen. Ist n gerade, hat h^{n-1} das Vorzeichen von h

$$\Rightarrow h^{n-1} f(\xi + \Theta h) > 0.$$

Nach (*) hat dann f in ξ ein strenges lokales Minimum.

□

25.14 Beispiele

(a) Für $f(x) = x^3$ gilt:

$$\begin{aligned} f'(x) &= 3x^2 & f'(0) &= 0 \\ f''(x) &= 6x & f''(0) &= 0 \\ f'''(x) &= 6 & f'''(0) &= 6 \neq 0. \end{aligned}$$

Also liegt kein Extremum in 0 vor.

(b) Für $f(x) = x^4$ gilt:

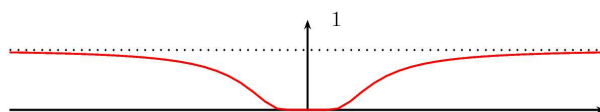
$$\begin{aligned} f'(x) &= 4x^3 & f'(0) &= 0 \\ f''(x) &= 12x^2 & f''(0) &= 0 \\ f'''(x) &= 24x & f'''(0) &= 0 \\ f^{(4)}(x) &= 24 & f^{(4)}(0) &> 0. \end{aligned}$$

Somit liegt in 0 ein strenges Minimum vor.

(c) Es gibt Fälle, in denen Satz 25.13 nicht anwendbar ist, z.B.

$$f(x) = \begin{cases} e^{-\frac{1}{x^2}} & (x \neq 0) \\ 0 & (x = 0) \end{cases}$$

Induktiv kann man zeigen, dass $f \in C^\infty(\mathbb{R})$ und $f^{(k)}(0) = 0 \ \forall k \in \mathbb{N}$.



Trotzdem liegt in 0 ein strenges lokales Minimum vor. Satz 25.13 liefert somit nur eine hinreichende Bedingung für lokale Extrema, die jedoch nicht notwendig ist.

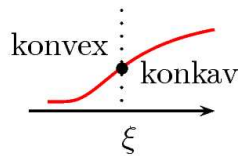
25.15 Definition: (Wendepunkt)

Eine in ξ differenzierbare Funktion $f(x)$ hat einen Wendepunkt in ξ , falls $f'(x)$ in ξ ein lokales Extremum hat.

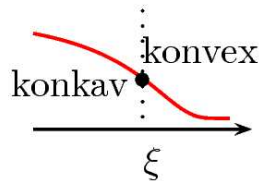
25.16 Bemerkung:

Im Wendepunkt wechselt das Krümmungsverhalten:

- Hat $f'(x)$ in ξ ein lokales Maximum, so ist f' in einer Umgebung links von ξ monoton wachsend und rechts von ξ monoton fallend.
(Konvex-konkav-Wechsel)



- Hat $f'(x)$ in ξ ein lokales Minimum, so liegt ein Konkav-konvex-Wechsel vor.



Analog zu Satz 25.11 auf Seite 187 und Satz 25.13 auf Seite 188 existieren daher notwendige und hinreichende Kriterien für Wendepunkte.

25.17 Satz: (Notwendige Bedingung für Wendepunkte)

Sei f in ξ zweimal differenzierbar und habe dort einen Wendepunkt. Dann ist

$$f''(\xi) = 0.$$

25.18 Satz: (Hinreichende Bedingung für Wendepunkte)

Sei $f \in C^{n+1}(I)$, $n \geq 2$, ξ innerer Punkt von I mit

$$f''(\xi) = f'''(\xi) = \dots = f^{(n)}(\xi) = 0 \text{ und } f^{(n+1)}(\xi) \neq 0.$$

- (a) *Ist n gerade, so liegt in ξ ein Wendepunkt vor.
Für $f^{(n+1)}(\xi) > 0$ ist es ein Konkav-konvex-Wechsel,
für $f^{(n+1)}(\xi) < 0$ ist es ein Konvex-konkav-Wechsel.*
- (b) *Falls n ungerade ist, handelt es sich um keinen Wendepunkt.*

25.19 Beispiel:

Für $f(x) = x^3$ gilt:

$$\begin{array}{ll} f'(x) = 3x^2 & f'(0) = 0 \\ f''(x) = 6x & f''(0) = 0 \\ f'''(x) = 6 & f'''(0) = 6 > 0. \end{array}$$

In $\xi = 0$ liegt daher ein Wendepunkt mit Konkav-konvex-Wechsel vor.

25.20 Kurvendiskussion

Ziel der Kurvendiskussion ist die Feststellung des qualitativen und quantitativen Verhaltens des Graphen einer Funktion $f(x)$. Hierzu gehören typischerweise folgende Untersuchungen:

- (a) Maximaler Definitionsbereich:

Beispiel: $f(x) = \frac{2x^2 + 3x - 4}{x^2}$ hat den Definitionsbereich $\mathbb{R} \setminus \{0\}$.

- (b) Symmetrien:

- (i) $f(x)$ ist symmetrisch zur y-Achse (gerade Funktion), falls

$$f(-x) = f(x) \quad \forall x$$

Beispiel: $f(x) = \cos x$ ist gerade Funktion.

- (ii) $f(x)$ ist symmetrisch zum Ursprung (ungerade Funktion), falls

$$f(-x) = -f(x) \quad \forall x$$

- (c) Polstellen:

Hat $f(x)$ die Form $f(x) = \frac{g(x)}{(x-\xi)^k}$ mit $g(x)$ stetig in ξ und $g(\xi) \neq 0$, so besitzt $f(x)$

- für ungerades k einen Pol mit Vorzeichenwechsel
- für gerades k einen Pol ohne Vorzeichenwechsel.

Beispiel: $f(x) = \frac{2x^2 + 3x - 4}{x^2}$ hat in $\xi = 0$ einen Pol ohne Vorzeichenwechsel.

$$\lim_{x \rightarrow 0^\pm} f(x) = -\infty.$$

(d) Verhalten im Unendlichen:

Bestimme $\lim_{x \rightarrow \infty} f(x)$, $\lim_{x \rightarrow -\infty} f(x)$, falls existent.

Untersuchung auf Asymptoten:

Eine Gerade $y = ax + b$ heißt Asymptote von $f(x)$ für $x \rightarrow \pm\infty$, falls $\lim_{x \rightarrow \pm\infty} (f(x) - (ax + b)) = 0$ gilt:

a und b werden bestimmt durch

$$a = \lim_{x \rightarrow \pm\infty} \frac{f(x)}{x}, \quad b = \lim_{x \rightarrow \pm\infty} (f(x) - ax)$$

Beispiel: $f(x) = \frac{2x^2 + 3x - 4}{x^2}$

$$\begin{aligned} a &= \lim_{x \rightarrow \pm\infty} \frac{2x^2 + 3x - 4}{x^3} = \lim_{x \rightarrow \pm\infty} \left(\frac{2}{x} + \frac{3}{x^2} - \frac{4}{x^3} \right) = 0 \\ b &= \lim_{x \rightarrow \pm\infty} f(x) = 2 \\ &\Rightarrow y = 2 \quad \text{Asymptote.} \end{aligned}$$

(e) Nullstellen:

Können bei Polynomen vom Grad ≤ 4 analytisch bestimmt werden. Ansonsten muss man „raten“ oder numerische Verfahren (z.B. Bisektionsverfahren 20.9 auf Seite 152, andere Verfahren folgen) verwenden.

Beispiel:

$$\begin{aligned} f(x) &= \frac{2x^2 + 3x - 4}{x^2} \\ 0 &= 2x^2 + 3x - 4 \\ \Rightarrow x_{1/2} &= \frac{-3 \pm \sqrt{9 + 32}}{4} = \frac{-3 \pm \sqrt{41}}{4} \\ x_1 &\approx -2,35 \\ x_2 &\approx 0,85 \end{aligned}$$

(f) Extrema, Monotonieintervalle:

Beispiel: $f(x) = \frac{2x^2 + 3x - 4}{x^2}$

$$\begin{aligned}f'(x) &= \frac{x^2(4x+3) - 2x(2x^2+3x-4)}{x^4} \\&= \frac{-3x^2+8x}{x^4} \\&= \frac{8-3x}{x^3} \\f'(x) &= 0 \Leftrightarrow x_3 = \frac{8}{3} \\f\left(\frac{8}{3}\right) &\approx 2,56 \\f''(x) &= \frac{x^3(-3) - 3x^2(8-3x)}{x^6} = \frac{6x^3-24x^2}{x^6} = \frac{6x-24}{x^4} \\f''\left(\frac{8}{3}\right) &< 0 \Rightarrow f \text{ hat in } x_3 = \frac{8}{3} \text{ lok. Max.}\end{aligned}$$

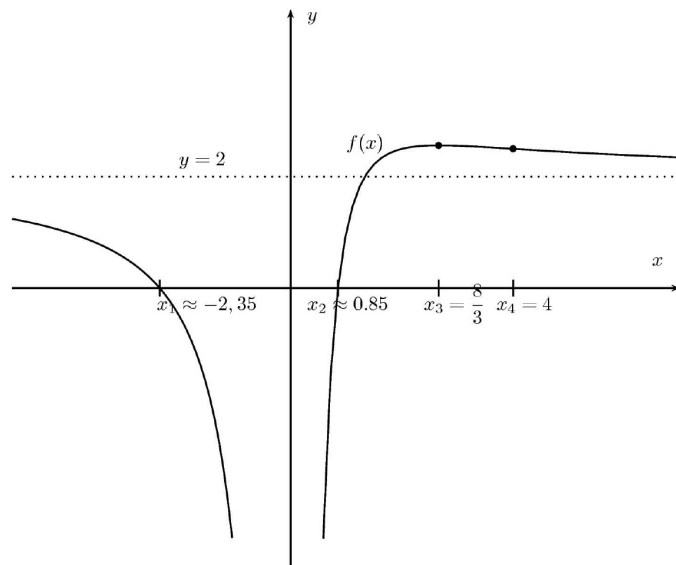
$$f'(x) = \begin{cases} < 0 & \text{für } \frac{8}{3} < x < \infty \text{ (streng mon. fallend)} \\ > 0 & \text{für } 0 < x < \frac{8}{3} \text{ (streng mon. wachsend)} \\ < 0 & \text{für } -\infty < x < 0 \text{ (streng mon. fallend)} \end{cases}$$

(g) Wendepunkte:

Beispiel: $f(x) = \frac{2x^2+3x-4}{x^2}$

$$\begin{aligned}f''(x) &= \frac{6x-24}{x^4} \\&\stackrel{!}{=} 0 \\&\Rightarrow x_4 = 4, \\f(x_4) &= \frac{5}{2} \\f'''(x) &= \frac{x^4 \cdot 6 - 4x^3(6x-24)}{x^8} = \frac{-18x^4+96x^3}{x^8} \\&= \frac{96-18x}{x^5} \\f'''(4) &> 0 : \quad \text{Konkav-konvex-Wechsel.}\end{aligned}$$

(h) Skizze:



Kapitel 26

Fixpunkt-Iteration

26.1 Motivation

Iterationsverfahren sind Verfahren der schrittweisen Annäherung. Sie gehören zu den wichtigsten Methoden in der Numerik zur Lösung von nichtlinearen Gleichungen, z.B. Bisektionsverfahren.

Verfahren zur iterativen Lösung einer nichtlinearen Gleichung $f(x) = 0$ mit einer C^1 -Funktion $f : D \rightarrow \mathbb{R}$, $D \subset \mathbb{R}$, haben meistens die Form

$$x_{k+1} = \Phi(x_k), \quad k = 0, 1, 2, \dots \quad (FPI)$$

Iterationen dieser Form heißen Fixpunkt-Iterationen mit Verfahrensfunktion Φ . Diese Bezeichnung ist wie folgt begründet:

Erzeugt die Iterationsvorschrift (FPI) eine konvergente Folge $(x_n)_{n \in \mathbb{N}_0}$ und ist Φ stetig, so folgt:

$$x^* := \lim_{k \rightarrow \infty} x_{k+1} = \lim_{k \rightarrow \infty} \Phi(x_k) = \Phi \left(\lim_{k \rightarrow \infty} (x_k) \right) = \Phi(x^*).$$

Das Verfahren konvergiert. Der Grenzwert x^* hat die Eigenschaft

$$x^* = \Phi(x^*)$$

und ist damit ein Fixpunkt der Verfahrensfunktion Φ .

Wie gewinnt man die Verfahrensfunktion Φ ?

Man formt die Gleichung $f(x) = 0$ in eine äquivalente Gleichung der Form $x = \Phi(x)$ um. Dies kann auf vielfältige Art geschehen und hat großen Einfluss auf das Konvergenzverhalten der Fixpunkt-Iteration.

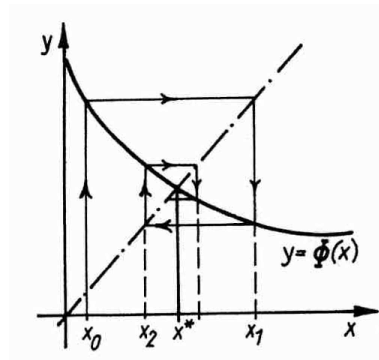


Abbildung 26.1: Fixpunkt-Iteration

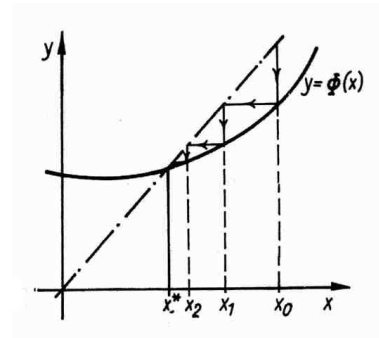
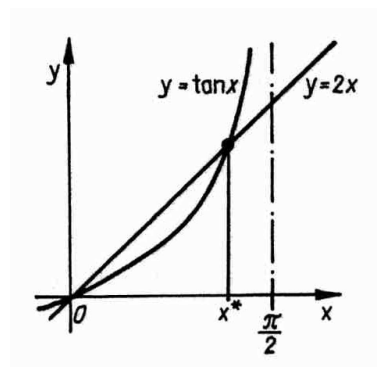


Abbildung 26.2: Fixpunkt-Iteration

26.2 Beispiel:

Suche (eindeutig bestimmte) Nullstelle $x^* \in \left(0, \frac{\pi}{2}\right)$ von $f(x) = 2x - \tan x = 0$



- (a) Die Umformung $2x - \tan x = 0 \Leftrightarrow x = \frac{1}{2} \tan x =: \Phi_a(x)$ führt mit der Anfangsnäherung, dem Startwert $x_0 := 1,2$, zu

k	x_k
0	1,2
1	1,286075811
2	1,708396214
3	-3,610761599
4	-0,253404236
$\vdots$	$\vdots$

Ab dieser Iteration konvergiert (x_k) monoton gegen 0.

- (b) Die Umformung $2x - \tan x = 0 \Leftrightarrow x = \arctan(2x) =: \Phi_b(x)$ führt mit gleichem Startwert $x_0 = 1,2$ zu

k	x_k
0	1,2
5	1,165657641
10	1,165561465
15	1,165561186
$\vdots$	$\vdots$

Man beobachtet Konvergenz dieser Iteration gegen die gesuchte nicht-triviale Nullstelle x^* .

Wie zeigt man Konvergenz einer Fixpunkt-Iteration?

Dazu zunächst:

26.3 Definition: (*Lipschitz-stetig, Lipschitz-Konstante*)

Eine Abbildung $\Phi : D \rightarrow \mathbb{R}$, $D \subset \mathbb{R}$, heißt Lipschitz-stetig auf D , falls es eine Konstante $L > 0$ gibt mit:

$$\forall x, y \in D : |\Phi(x) - \Phi(y)| \leq L \cdot |x - y|.$$

L heißt Lipschitz-Konstante.

Kann man $L < 1$ wählen, so heißt die Abbildung Φ kontrahierend und L nennt man dann auch Kontraktionskonstante von Φ .

26.4 Bemerkungen:

- (a) Φ Lipschitz-stetig $\Rightarrow \Phi$ stetig, denn: Aus $x_k \rightarrow x$ ($k \rightarrow \infty$) folgt

$$|\Phi(x_k) - \Phi(x)| \leq L \cdot |x_k - x| \rightarrow 0 \quad (k \rightarrow \infty)$$

d.h. $\Phi(x_k) \rightarrow \Phi(x)$ ($k \rightarrow \infty$).

- (b) Man beachte, dass die Kontraktionseigenschaft nur durch eine Abschätzung

$$|\Phi(x) - \Phi(y)| \leq L \cdot |x - y|, \quad L < 1 \text{ gesichert ist.}$$

Die schwächere Abschätzung

$$|\Phi(x) - \Phi(y)| \leq L \cdot |x - y| \quad \forall x, y \text{ mit } x \neq y \text{ leistet das nicht.}$$

26.5 Satz: (*Lipschitz-Stetigkeit von C^1 -Funktionen*)

Jede C^1 -Funktion $\Phi : [a, b] \rightarrow \mathbb{R}$ ist Lipschitz-stetig auf $[a, b]$ mit einer Lipschitz-Konstanten

$$L := \sup\{|\Phi'(x)| : a \leq x \leq b\}.$$

Ist $L < 1$, so ist Φ kontrahierend.

Ist $L > 1$, so ist Φ nicht kontrahierend.

Beweis:

Behauptung folgt direkt aus MWS:

$$|\Phi(x) - \Phi(y)| = \Phi'(\xi) \cdot |x - y| \leq L \cdot |x - y| \quad \text{mit } \xi \in (x, y).$$

Von zentraler Bedeutung ist

26.6 Satz: (Banachscher Fixpunktsatz, 1D-Version)

Es sei $D \subset \mathbb{R}$ abgeschlossen und $\Phi : D \rightarrow D$ eine kontrahierende Abbildung von D in sich (d.h. $f(D) \subset D$) mit Kontraktionskonstante $L < 1$. Dann gelten die folgenden Aussagen:

- (a) Es gibt genau einen Fixpunkt x^* von Φ in D .
- (b) Für jeden Startwert $x_0 \in D$ konvergiert die Fixpunkt-Iteration $x_{k+1} = \Phi(x_k)$ gegen den Fixpunkt x^* .
- (c) Es gelten die Fehlerabschätzungen (a posteriori und a priori Abschätzungen):

$$|x_n - x^*| \leq \frac{L}{1-L} |x_n - x_{n-1}| \leq \frac{L^n}{1-L} |x_1 - x_0|.$$

Beweis:

Sei beliebiges $x_0 \in D$ gegeben. Da $\Phi : D \rightarrow D$, ist die Folge $(x_k)_{k \in \mathbb{N}_0}$ durch $x_{k+1} = \Phi(x_k)$ eindeutig definiert und für $k \in \mathbb{N}$ gilt:

$$\begin{aligned} |x_{k+1} - x_k| &= |\Phi(x_k) - \Phi(x_{k-1})| \\ &\leq L |x_k - x_{k-1}| \\ &\leq L^2 |x_{k-1} - x_{k-2}| \leq \dots \quad (*) \end{aligned}$$

Durch Iteration der (ersten) Abschätzung folgt für $k \geq n$:

$$|x_{k+1} - x_k| \leq L^{k+1-n} |x_n - x_{n-1}| \quad (**)$$

und damit auch für $m \geq n$:

$$\begin{aligned}
 |x_m - x_n| &= |(x_m - x_{m-1}) + (x_{m-1} - x_{m-2}) + (x_{m-2} - x_{m-3}) + \dots + (x_{n+1} - x_n)| \\
 &\leq \sum_{k=n}^{m-1} |x_{k+1} - x_k| \quad (\Delta\text{-Ungleichung}) \\
 &\leq \left(\sum_{k=n}^{m-1} L^{k+1-n} \right) |x_n - x_{n-1}| \\
 &\leq \left(\sum_{j=1}^{\infty} L^j \right) |x_n - x_{n-1}| \\
 &= L \cdot \left(\sum_{j=0}^{\infty} L^j \right) |x_n - x_{n-1}| \\
 &= \frac{L}{1-L} |x_n - x_{n-1}| \\
 &\leq \frac{L^n}{1-L} |x_1 - x_0| \quad (\text{vgl. (*) oder (**)}).
 \end{aligned}$$

Also gilt für $m \geq n$:

$$|x_m - x_n| \leq \frac{L}{1-L} |x_n - x_{n-1}| \leq \frac{L^n}{1-L} |x_1 - x_0| \quad (***)$$

Wegen $L < 1$ gilt $L^n \rightarrow 0$ ($n \rightarrow \infty$) und damit ist $(x_n)_{n \in \mathbb{N}_0}$ eine Cauchy-Folge, also auch konvergent mit einem Grenzwert x^* . Da D abgeschlossen ist (**), gilt $x^* \in D$.

Φ ist stetig (vgl. 26.3 (a)) und damit ist der Grenzwert x^* Fixpunkt von Φ . x^* ist der einzige Fixpunkt in D . Denn wäre $x^{**} \neq x^*$ ein weiterer Fixpunkt, so würde folgen:

$$|x^{**} - x^*| = |\Phi(x^{**}) - \Phi(x^*)| \leq L|x^{**} - x^*| < |x^{**} - x^*|, \text{ ein Widerspruch.}$$

Die behauptete „Fehlerabschätzung“ folgen direkt aus (**), wenn $m \rightarrow \infty$ strebt.

□

26.7 Bemerkung:

Die Bedingungen des Satzes 26.6 auf der vorherigen Seite sind hinreichend, aber nicht notwendig für die Existenz und Eindeutigkeit eines Fixpunktes, sowie für die Konvergenz der Fixpunkt-Iteration.

26.8 Beispiel:

Man berechne den kleinsten Fixpunkt von $\Phi(x) = 0,1 \cdot e^x$. Wir setzen $D := [-1, 1]$.

Für $x \in D$ gilt: $0 < \Phi(x) \leq \frac{e}{10} < 1$, also $\Phi : [-1, 1] \rightarrow [-1, 1]$. Es gilt $\Phi'(x) = \Phi(x)$. Also ist Φ kontrahierend mit

$$L = \sup\{|\Phi'(x)| : x \in [-1, 1]\} = \frac{e}{10}.$$

Die Voraussetzungen des Fixpunktsatzes sind erfüllt. x^* soll mit einem absoluten Fehler kleiner oder gleich 10^{-6} berechnet werden. Wir starten mit $x_0 := 1$ und $x_1 = 0,2718281828$. Die Forderung

$$|x^* - x_n| \leq \frac{L^n}{1 - L} |x_1 - x_0| \leq 10^{-6}$$

führt mit den vorhandenen Werten auf

$$n \geq \frac{-6}{\log_{10} L} = 10,606\dots$$

Es genügen also $n = 11$ Iterationen, um die verlangte Genauigkeit zu garantieren. Tatsächlich ist nach 11 Iterationen bereits eine zehnstellige Genauigkeit erreicht: $x^* = 0,1118325592$.

Kapitel 27

Das Newton-Verfahren

27.1 Motivation

Das Konvergenzverhalten von Fixpunkt-Iterationen zur Lösung von nichtlinearen Gleichungen hängt entscheidend von der Wahl der Verfahrensfunktion Φ ab.

Eine geschickte Wahl von Φ führt auf das Newton-Verfahren. Es beruht auf einer Taylorapproximation 1. Ordnung (Linearisierung) und konvergiert schnell, falls es konvergiert.

27.2 Newton-Verfahren für eine nichtlineare Gleichung

Wir suchen eine Nullstelle einer skalaren C^1 -Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$, d.h. eine Lösung der Gleichung $f(x) = 0$.

Das Newton-Verfahren besteht darin, bei einem Näherungswert x_0 den Graphen von f durch die Tangente zu ersetzen und dessen Nullstelle als neue Näherung x_1 zu benutzen. Dieses Vorgehen wird iteriert.

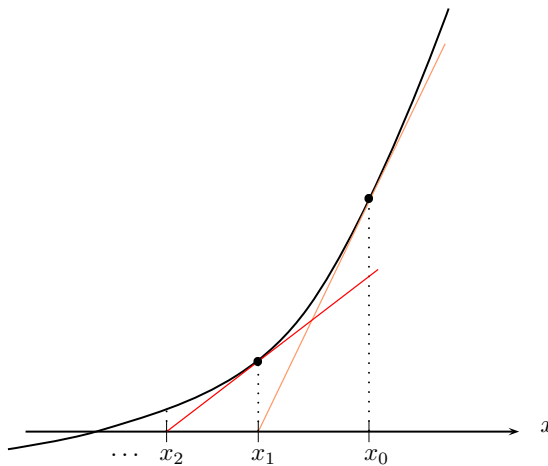
Tangente in x_0 :

$$T_1(x) = f(x_0) + (x - x_0) \cdot f'(x_0)$$

Taylorpolynom 1. Ordnung

x_1 := Nullstelle von T_1 :

$$\begin{aligned} 0 &\stackrel{!}{=} f(x_0) + (x_1 - x_0) \cdot f'(x_0) \\ \Rightarrow x_1 &= x_0 - \frac{f(x_0)}{f'(x_0)} \end{aligned}$$



allgemeine Iterationsvorschrift:

$$x_{k+1} = x_k - \frac{f(x_k)}{f'(x_k)} =: \Phi(x_k), \quad k = 0, 1, 2, \dots$$

27.3 Beispiele

$f(x) = x^2 - 2$ hat Nullstelle bei $\pm\sqrt{2} \approx \pm 1,412136\dots$
Mit der Iterationsvorschrift

$$x_{k+1} = x_k - \frac{x_k^2 - 2}{2x_k}$$

und Startwert $x_0 := 1$ ergibt sich:

k	x_k
0	1,0
1	1.5000000
2	1.4166667
3	1.4142157
4	1.4142136

27.4 Bemerkung:

Das Verfahren konvergiert nicht immer, im allgemeinen konvergiert es erst, wenn der Startwert x_0 „hinreichend nahe“ bei der Nullstelle liegt (lokale Konvergenz).

Einen wichtigen Fall, in dem Konvergenz auftritt, enthält der folgende Satz:

27.5 Satz: (Konvergenz des Newtonverfahrens)

Es sei $f : [a, b] \rightarrow \mathbb{R}$ eine zweimal differenzierbare konvexe Funktion mit $f(a) < 0$ und $f(b) > 0$. Dann gilt:

- (a) Es gibt genau ein $\xi \in (a, b)$ mit $f(\xi) = 0$.
- (b) Ist $x_0 \in [a, b]$ beliebig mit $f(x_0) \geq 0$, so ist die Folge

$$x_{n+1} := x_n - \frac{f(x_n)}{f'(x_n)}, \quad n \in \mathbb{N},$$

wohldefiniert und konvergiert monoton fallend gegen ξ .

- (c) Gilt $f'(\xi) \geq C > 0$ und $f''(x) \leq K \quad \forall x \in (\xi, b)$, so hat man für $n \in \mathbb{N}$ die Abschätzung

$$|x_{n+1} - x_n| \leq |\xi - x_n| \leq \frac{K}{2C} |x_n - x_{n-1}|^2.$$

Beweis:

von (a) und (b):

- (a) Da f konvex ist, gilt $f''(x) \geq 0 \quad \forall x \in (a, b)$ und folglich ist f' monoton wachsend in $[a, b]$. Als stetige Funktion nimmt f sein Minimum über dem Intervall $[a, b]$ an, genauer:

$$\exists q \in [a, b] \text{ mit } f(q) = \inf\{f(x) : x \in [a, b]\} < 0.$$

Falls $q \neq a$, gilt $f'(q) = 0$ und damit $f'(x) \leq 0$ für $x \leq q$ (f' monoton wachsend).

$\Rightarrow f$ ist im Intervall $[a, q]$ monoton fallend.

$\Rightarrow f$ kann in $[a, q]$ keine Nullstelle haben.

Dies stimmt auch für $q = a$.

Nach dem Zwischenwertsatz 20.9 auf Seite 152 (b) liegt im Intervall (q, b) mindestens eine Nullstelle von f . Nach dem bisher Gezeigten liegen alle Nullstellen in (q, b) .

Angenommen es gäbe zwei Nullstellen $\xi_1 < \xi_2$.

Nach dem Mittelwertsatz 23.2 auf Seite 173 existiert $t \in (q, \xi_1)$ mit

$$f'(t) = \frac{f(\xi_1) - f(q)}{\xi_1 - q} = \frac{-f(q)}{\xi_1 - q} > 0.$$

$\Rightarrow f'(x) > 0$ für alle $x \geq \xi_1$ (f' monoton wachsend)

$\Rightarrow f$ ist streng monoton wachsend (und damit positiv) in $(\xi_1, b]$ und kann damit keine zweite Nullstelle $\xi_2 \in (\xi_1, b]$ besitzen.

- (b) Sei $x_0 \in [a, b]$ mit $f(x_0) \geq 0$. Nach bisher Gezeigtem gilt $x_0 \geq \xi$. Durch Induktion beweisen wir, dass für die Folge $(x_n)_{n \in \mathbb{N}_0}$ definiert durch

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

gilt:

$$f(x_n) \geq 0 \text{ und } x_{n-1} \geq x_n \geq \xi \quad \forall n \in \mathbb{N}_0.$$

Ind. Anfang: $n=0$.

$$n \rightarrow n+1: \text{ Aus } x_n \geq \xi \text{ folgt } f'(x_n) \geq f'(\xi) > 0 \quad (\text{vgl. (a)})$$

$$\Rightarrow \frac{f(x_n)}{f'(x_n)} \geq 0 \stackrel{(*)}{\Rightarrow} x_{n+1} \leq x_n.$$

Wir zeigen: $f(x_{n+1}) \geq 0$.

Wir betrachten dazu die Hilfsfunktion

$$g(x) := f(x) - f(x_n) - (x - x_n)f'(x_n).$$

f ist monoton wachsend.

$$\Rightarrow g'(x) = f'(x) - f'(x_n) \leq 0 \text{ für } x \leq x_n.$$

Da $g(x_n) = 0$, gilt $g(x) \geq 0$ für $x \leq x_n$, also gilt insbesondere

$$\begin{aligned} 0 \leq g(x_{n+1}) &= f(x_{n+1}) - \underbrace{f(x_n) - (x_{n+1} - x_n)f'(x_n)}_{= 0 \text{ nach } (*)} \\ &= f(x_{n+1}) \end{aligned}$$

Aus $f(x_{n+1}) \geq 0$ folgt nach (a), dass $x_{n+1} \geq \xi$. Damit ist gezeigt: $(x_n)_{n \in \mathbb{N}_0}$ fällt monoton und ist nach unten beschränkt durch ξ .

$\Rightarrow$ Es existiert $\lim_{n \rightarrow \infty} x_n =: x^*$. Man hat $f(x^*) = 0$ (Fixpunkt-Iteration).

$\Rightarrow x^* = \xi$ wegen der Eindeutigkeit der Nullstelle.

□

(c) ohne Beweis.

27.6 Bemerkung:

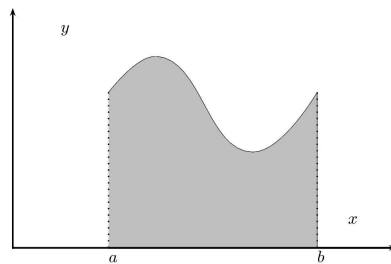
- (a) Analoge Aussagen gelten auch, falls f konkav ist oder $f(a) > 0$ und $f(b) < 0$ gilt.
- (b) Die Fehlerabschätzung 27.5 auf der vorherigen Seite (c) besagt, dass beim Newton-Verfahren quadratische Konvergenz vorliegt. Ist etwa $\frac{K}{2C} \approx 1$ und stimmen x_{n-1} und x_n auf K Dezimalen überein, so ist die Näherung x_{n+1} auf $2k$ Dezimalen genau und bei jeder Iteration verdoppelt sich die Zahl der gültigen Stellen.

Kapitel 28

Das Bestimmte Integral

28.1 Motivation

Es sei $f(x), x \in [a, b]$ eine reellwertige Funktion.



Ziel: Berechnung der Fläche zwischen $f(x)$ und der x -Achse.

Annahme: $f : [a, b] \rightarrow \mathbb{R}$ sei beschränkt.

28.2 Definition: (Zerlegung, Partition, Feinheit)

(a) Eine Menge der Form

$$Z = \{a = x_0 < x_1 < \dots < x_n = b\}$$

heißt Zerlegung (Partition, Unterteilung) des Intervalls $[a, b]$.
Die x_i heißen Knoten der Zerlegung.

(b) $\|Z\| := \max_{0 \leq i \leq n-1} |x_{i+1} - x_i|$ heißt Feinheit der Zerlegung Z .

- (c) $\mathfrak{Z} := \mathfrak{Z}[a, b] : \text{Menge aller Zerlegungen des Intervalls } [a, b].$
- (d) Eine Zerlegung Z_1 ist eine feinere Zerlegung als Z_2 ($Z_1 \supset Z_2$, Verfeinerung), falls Z_1 durch Hinzunahme weiterer Knoten zu Z_2 entsteht.

28.3 Definition: (Riemann-Summe)

- Jede Summe der Form

$$R_f(Z) := \sum_{i=0}^{n-1} f(\xi_i) \cdot (x_{i+1} - x_i), \quad \xi_i \in [x_i, x_{i+1}]$$

heißt Riemann-Summe von f zur Zerlegung Z .

-

$$U_f(Z) := \sum_{i=0}^{n-1} \left(\inf_{\xi \in [x_i, x_{i+1}]} f(\xi) \right) \cdot (x_{i+1} - x_i)$$

heißt Untersumme von f zur Zerlegung Z .

-

$$O_f(Z) := \sum_{i=0}^{n-1} \left(\sup_{\xi \in [x_i, x_{i+1}]} f(\xi) \right) \cdot (x_{i+1} - x_i)$$

heißt Obersumme von f zur Zerlegung Z .

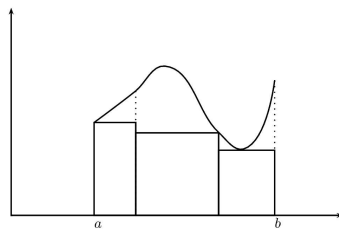


Abbildung 28.1:

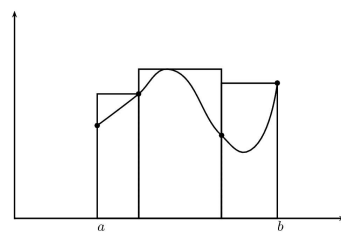


Abbildung 28.2:

28.4 Bemerkung:

- (a) Für jedes feste Z gilt: $U_f(Z) \leq R_f(Z) \leq O_f(Z)$.
- (b) Verfeinerungen vergrößern Untersummen und verkleinern Obersummen:

$$Z_1 \supset Z_2 \Rightarrow U_f(Z_1) \geq U_f(Z_2), \quad O_f(Z_1) \leq O_f(Z_2).$$

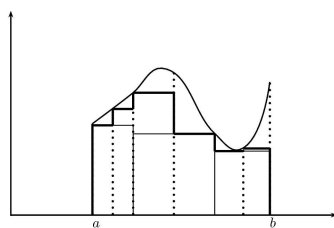


Abbildung 28.3: $U_f(Z_1) \leq U_f(Z_2)$

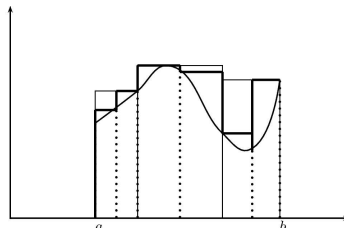


Abbildung 28.4: $O_f(Z_1) \geq O_f(Z_2)$

(c) Für zwei beliebige Zerlegungen Z_1, Z_2 von $[a, b]$ gilt stets

$$U_f(Z_1) \leq O_f(Z_2).$$

Insbesondere sind $\left\{ \begin{array}{c} \text{Untersummen} \\ \text{Obersummen} \end{array} \right\}$ nach $\left\{ \begin{array}{c} \text{oben} \\ \text{unten} \end{array} \right\}$ beschränkt.

28.5 Definition: (Riemann-Integral)

- Aufgrund der obigen Eigenschaften existieren die Grenzwerte

$$\int_a^b f(x) \, dx := \sup_{Z \in \mathfrak{Z}[a,b]} U_f(Z) \quad \text{Riemann'sches Unterintegral}$$

$$\int_a^b f(x) \, dx := \inf_{Z \in \mathfrak{Z}[a,b]} O_f(Z) \quad \text{Riemann'sches Oberintegral.}$$

- $f(x)$ heißt (Riemann-)integrierbar über $[a, b]$, wenn Ober- und Unterintegral übereinstimmen.

Dann heißt

$$\int_a^b f(x) \, dx := \int_a^b f(x) \, dx = \int_a^b f(x) \, dx$$

das (Riemann-) Integral von $f(x)$ über $[a, b]$.

28.6 Beispiele

(a) Es sei $f(x) = c = \text{const}$ auf $[a, b]$.

$$\Rightarrow U_f(Z) = O_f(Z) = \sum_{i=0}^{n-1} c \cdot (x_{i+1} - x_i) = c \cdot (b - a).$$

D.h. $f(x)$ ist integrierbar mit $\int_a^b f(x) \, dx = c \cdot (b - a)$.

(b) Es sei $\xi \in [a, b]$. Dann ist

$$f(x) = \begin{cases} 0 & (x \neq \xi) \\ 1 & (x = \xi) \end{cases} \quad \text{integrierbar.}$$

Denn für jede Zerlegung Z gilt:

$$U_f(Z) = 0, \quad 0 < O_f(Z) \leq 2 \cdot \|Z\|.$$

Also ist $\int_a^b f(x) \, dx = 0$ (da $\|Z\|$ beliebig klein gewählt werden kann).

(c) Es sei

$$f(x) = \begin{cases} 0 & (x \in [0, 1] \cap \mathbb{Q}) \\ 1 & (x \in [0, 1] \setminus \mathbb{Q}) \end{cases} \quad (\text{Dirichlet'sche Sprungfunktion}).$$

Für jede Zerlegung Z ist

$$U_f(Z) = 0, O_f(Z) = 1.$$

Somit ist $f(x)$ nicht Riemann-integrierbar (Es gibt jedoch allgemeine Integrierbarkeitsbegriffe - z.B. Lebesgue-Integral -, nach denen auch solche Funktionen integriert werden können).

28.7 Definition: (Eigenschaften des Riemann-Integrals)

Ist $f : [a, b] \rightarrow \mathbb{R}$ eine nichtnegative, Riemann-integrierbare Funktion, so wird die Zahl

$$\int_a^b f(x) \, dx$$

als Fläche zwischen der Kurve $f(x)$ und der x -Achse zwischen den Geraden $x = a$ und $x = b$ bezeichnet.

28.8 Bemerkung:

- Ist f negativ, so kann man die Fläche durch $\int_a^b |f(x)| \, dx$ definieren.
- Ist $b < a$, so definiert man $\int_a^b f(x) \, dx := - \int_b^a f(x) \, dx$.

28.9 Lemma: (*Monotonie des Riemann-Integrals*)

Sind f und g integrierbar über $[a, b]$ mit $f(x) \leq g(x) \quad \forall x \in [a, b]$, so gilt:

$$\int_a^b f(x) \, dx \leq \int_a^b g(x) \, dx.$$

Beweis:

Zu jeder vorgegebenen Zerlegung Z gilt $U_f(Z) \leq U_g(Z)$, $O_f(Z) \leq O_g(Z)$. Diese Ungleichungen übertragen sich auf Supremum (Unterintegral) bzw. Infimum (Oberintegral).

□

28.10 Folgerung:

$$\left| \int_a^b f(x) \, dx \right| \leq \int_a^b |f(x)| \, dx$$

Beweis:

Nach Lemma 28.9 folgt aus $-|f(x)| \leq f(x) \leq |f(x)|$ die Beziehung

$$-\int |f(x)| \, dx \leq \int f(x) \, dx \leq \int |f(x)| \, dx$$

und somit die Behauptung.

□

28.11 Folgerung (Nichtnegativität):

Integration erhält Nichtnegativität:

$$f(x) \geq 0 \quad \forall x \in [a, b] \Rightarrow \int_a^b f(x) \, dx \geq 0.$$

Beweis:

Lemma 28.9 angewendet auf $f(x)$ und 0 ergibt

$$\int_a^b f(x) \, dx \geq \int_a^b 0 \, dx = 0.$$

28.12 Satz: (Linearität der Integration)

Es seien $f(x), g(x)$ integrierbar über $[a, b]$ so wie $\alpha, \beta \in \mathbb{R}$. Dann ist auch $\alpha \cdot f(x) + \beta \cdot g(x)$ integrierbar, und es gilt

$$\int_a^b (\alpha \cdot f(x) + \beta \cdot g(x)) \, dx = \alpha \cdot \int_a^b f(x) \, dx + \beta \cdot \int_a^b g(x) \, dx.$$

Beweis:

Aufgrund der Ungleichungen

$$\begin{aligned} \sup_{x \in [x_i, x_{i+1}]} (f(x) + g(x)) &\leq \sup_{x \in [x_i, x_{i+1}]} f(x) + \sup_{x \in [x_i, x_{i+1}]} g(x) \\ \inf_{x \in [x_i, x_{i+1}]} (f(x) + g(x)) &\geq \inf_{x \in [x_i, x_{i+1}]} f(x) + \inf_{x \in [x_i, x_{i+1}]} g(x) \end{aligned}$$

gilt für jede Zerlegung Z von $[a, b]$

$$\begin{aligned} U_{f+g}(Z) &\geq U_f(Z) + U_g(Z), \\ O_{f+g}(Z) &\leq O_f(Z) + O_g(Z) \end{aligned}$$

also

$$\begin{aligned} \int_a^b (f(x) + g(x)) \, dx &\geq \int_a^b f(x) \, dx + \int_a^b g(x) \, dx, \\ \int_a^b (f(x) + g(x)) \, dx &\leq \int_a^b f(x) \, dx + \int_a^b g(x) \, dx. \end{aligned}$$

Daraus folgt die Additivität.

Für $\alpha \geq 0$ und jede Zerlegung Z von $[a, b]$ gilt weiter

$$U_{\alpha f}(Z) = \alpha U_f(Z), \quad O_{\alpha f}(Z) = \alpha O_f(Z),$$

damit

$$\begin{aligned} \int_a^b \alpha \cdot f(x) \, dx &= \alpha \cdot \int_a^b f(x) \, dx, \\ \int_a^b \alpha \cdot f(x) \, dx &= \alpha \cdot \int_a^b f(x) \, dx. \end{aligned}$$

Für $\alpha < 0$ hat man stattdessen

$$U_{\alpha f}(Z) = \alpha O_f(Z), \quad O_{\alpha f}(Z) = \alpha U_f(Z)$$

und somit

$$\begin{aligned} \int_a^b \alpha \cdot f(x) \, dx &= \alpha \cdot \int_a^b f(x) \, dx, \\ \int_a^b \alpha \cdot f(x) \, dx &= \alpha \cdot \int_a^b f(x) \, dx. \end{aligned}$$

In beiden Fällen folgt $\int_a^b \alpha \cdot f(x) \, dx = \alpha \cdot \int_a^b f(x) \, dx$.

□

28.13 Lemma: (Zusammensetzung von Integrationsintervallen)

Ist $a \leq c \leq b$, so ist f genau dann über $[a, b]$ integrierbar, wenn f über $[a, c]$ und über $[c, b]$ integrierbar ist. Dann gilt:

$$\int_a^b f(x) \, dx = \int_a^c f(x) \, dx + \int_c^b f(x) \, dx.$$

Beweis:

Aussage folgt unmittelbar aus der Betrachtung von Ober- und Untersummen zu Zerlegungen von $[a, b]$, die c als Knoten enthalten. (Beliebige Zerlegungen von $[a, b]$ können ggf. durch Hinzunahme von c verfeinert werden).

□

Kriterien für Integrierbarkeit

28.14 Lemma: (Riemann'sches Kriterium)

Ist $f : [a, b] \rightarrow \mathbb{R}$ beschränkt, so sind die folgenden Aussagen äquivalent:

- (i) f ist integrierbar über $[a, b]$.
- (ii) Für jedes $\epsilon > 0$ existiert eine Zerlegung Z von $[a, b]$ derart, dass gilt:

$$O_f(Z) - U_f(Z) < \epsilon.$$

Beweis:

„(i) $\Rightarrow$ (ii)“: Nach Definition des Ober- und Unterintegrals gibt es zu jedem $\epsilon > 0$ Zerlegungen Z_1, Z_2 mit

$$\begin{aligned} 0 &\leq O_f(Z_1) - \int_a^b f(x) \, dx < \frac{\epsilon}{2}, \\ 0 &\leq \int_a^b f(x) \, dx - U_f(Z_2) < \frac{\epsilon}{2}. \end{aligned}$$

Diese Ungleichung bleiben beim Übergang zu Verfeinerungen gültig, also kann $Z = Z_1 = Z_2$ gewählt werden. Addition beider Ungleichungen ergibt dann (ii).

„(ii) $\Rightarrow$ (i)“: Für jedes $\epsilon > 0$ folgt aus (ii), dass

$$0 \leq \int_a^b f(x) \, dx - \int_a^b f(x) \, dx \leq O_f(Z) - U_f(Z) < \epsilon.$$

Da aber Ober- und Unterintegral identisch sind, ist f integrierbar über $[a, b]$.

□

Wofür kann man das Riemann'sche Kriterium anwenden?

28.15 Satz: (Integrierbarkeit monotoner Funktionen)

Sei $f : [a, b] \rightarrow \mathbb{R}$ beschränkt.

Ist f monoton, so ist f integrierbar.

Beweis:

Sei o.B.d.A. f monoton wachsend. Betrachte die äquidistante Zerlegung Z mit

$x_i = a + \frac{i}{n}(b - a)$ und $i = 0, \dots, n$.

Dann gilt:

$$\begin{aligned} O_f(Z) - U_f(Z) &= \sum_{i=0}^{n-1} (f(x_{i+1}) - f(x_i)) \cdot (x_{i+1} - x_i) && (\text{da } f \text{ wachsend}) \\ &= \frac{b-a}{n} \sum_{i=0}^{n-1} (f(x_{i+1}) - f(x_i)) && (\text{Äquidist. Zerlegung}) \\ &= \frac{b-a}{n} \underbrace{(f(b) - f(a))}_{\text{beschränkt}} && (\text{„Ziehharmonikasumme“}) \end{aligned}$$

Für $n \rightarrow \infty$ konvergiert dieser Ausdruck gegen 0. Damit ist f nach dem Riemann'schen Kriterium integrierbar.

□

Gibt es noch eine weitere Anwendung?

28.16 Satz: (Integrierbarkeit stetiger Funktionen)

Sei $f : [a, b] \rightarrow \mathbb{R}$ beschränkt.

Ist f stetig, so ist f integrierbar.

Beweis:

Nach Satz 20.10 auf Seite 154 ist die stetige Funktion f auf dem abgeschlossenen Intervall $[a, b]$ gleichmäßig stetig. Seien $\epsilon, \delta > 0$ so gewählt, dass:

$$|x - y| < \delta \quad \Rightarrow \quad |f(x) - f(y)| < \frac{\epsilon}{b-a} \quad (*)$$

Für eine Zerlegung Z mit $\|Z\| < \delta$ folgt dann:

$$\begin{aligned} O_f(Z) - U_f(Z) &= \sum_{i=0}^{n-1} \left(\sup_{x \in [x_i, x_{i+1}]} f(x) - \inf_{x \in [x_i, x_{i+1}]} f(x) \right) \cdot (x_{i+1} - x_i) \\ &\leq \sum_{i=0}^{n-1} \frac{\epsilon}{b-a} (x_{i+1} - x_i) \quad (\text{nach } (*)) \\ &= \frac{\epsilon}{b-a} (a-b) = \epsilon \quad (\text{„Ziehharmonikasumme“}) \end{aligned}$$

Damit ist f nach dem Riemann'schen Kriterium integrierbar.

□

Welche Konsequenzen hat dieser Satz?

28.17 Beispiele

Folgende Funktionen sind über einem Intervall $[a, b]$ integrierbar.

- (a) $f(x) = \sqrt{x}$ falls $a \geq 0$.
- (b) $f(x) = e^x$
- (c) $f(x) = \ln x$ falls $a > 0$.
- (d) $f(x) = \sin x$
- (e) Polynomfunktionen:

$$f(x) = \sum_{k=0}^n a_k x^k.$$

- (f) Stückweise stetige Funktionen mit endlich vielen Sprungstellen
(Induktionsbeweis über die Zahl der Sprungstellen).

Welche Eigenschaften kann man für Integrale mit stetigen Funktionen noch zeigen?

28.18 Satz: (Mittelwertsatz der Integralrechnung)

Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig und $g : [a, b] \rightarrow \mathbb{R}$ stückweise stetig und nichtnegativ.
Dann existiert ein $\xi \in [a, b]$ mit

$$\int_a^b f(x)g(x) \, dx = f(\xi) \int_a^b g(x) \, dx.$$

Bemerkung: Für $\int_a^b g(x) \, dx > 0$ kann $f(\xi) = \frac{\int_a^b f(x)g(x) \, dx}{\int_a^b g(x) \, dx}$ als Mittelwert von f im Intervall $[a, b]$ mit einer Gewichtsfunktion g angesehen werden.

Beweis:

(von Satz 28.18 auf der vorherigen Seite)

Seien $m := \min_{x \in [a, b]} f(x)$, $M := \max_{x \in [a, b]} f(x)$. Da g nichtnegativ ist, gilt:

$$m \cdot g(x) \leq f(x) \cdot g(x) \leq M \cdot g(x) \quad \forall x \in [a, b]$$

und damit nach der Monotonieregel (Lemma 28.9 auf Seite 209)

$$m \cdot \int_a^b g(x) \, dx \leq \int_a^b f(x)g(x) \, dx \leq M \cdot \int_a^b g(x) \, dx$$

Somit existiert eine Zahl $\mu \in [m, M]$ („Mittelwert“) mit

$$\mu \cdot \int_a^b g(x) \, dx = \int_a^b f(x)g(x) \, dx$$

Nach dem Zwischenwertsatz 20.9 (b) existiert dann ein $\xi \in [a, b]$ mit $f(\xi) = \mu$.
Somit gilt:

$$f(\xi) \cdot \int_a^b g(x) \, dx = \int_a^b f(x)g(x) \, dx.$$

□

Mit $g(x) = 1 \quad \forall x \in [a, b]$ ergibt sich eine wichtige Folgerung, die oft auch schon als Mittelwertsatz der Integralrechnung bezeichnet wird.

28.19 Korollar: (Integralmittel)

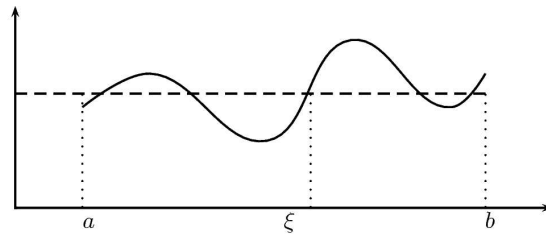
Für eine stetige Funktion $f : [a, b] \rightarrow \mathbb{R}$ gibt es ein $\xi \in [a, b]$ mit

$$f(\xi) = \frac{1}{b-a} \int_a^b f(x) \, dx.$$

28.20 Veranschaulichung

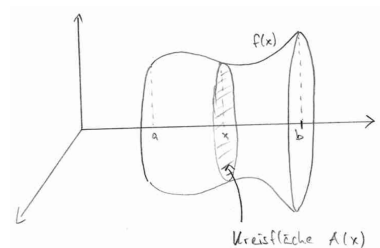
Die Fläche unter der Kurve (d.h. das Integral $\int_a^b f(x) \, dx$) stimmt mit der Rechtecksfläche $(b-a)f(\xi)$ überein:

$$\int_a^b f(x) \, dx = f(\xi)(b-a).$$



28.21 Volumen von Rotationskörpern

Eine Anwendung des bestimmten Integrals ist die Volumenberechnung von Rotationskörpern.



Kreisfläche $A(x)$ hat Radius $f(x)$

$$\Rightarrow A(x) = \pi (f(x))^2$$

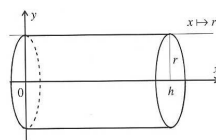
Prinzip von Cavalieri

Das Volumen des Rotationskörpers ergibt sich durch Integration der Flächeninhalte $A(x)$ im x -Intervall $[a, b]$

$$V = \int_a^b A(x) \, dx = \pi \int_a^b (f(x))^2 \, dx$$

28.22 Beispiel:

(a) Kreiszylinder:

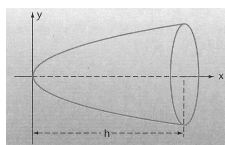


$$f(x) = r = \text{konstant.}$$

Volumen eines Kreiszylinders mit Radius r und Höhe h :

$$\begin{aligned} V &= \pi \int_0^h (f(x))^2 \, dx \\ &= \pi \int_0^h r^2 \, dx \\ &= \pi r^2 \int_0^h dx = \pi \cdot r^2 \cdot h \end{aligned}$$

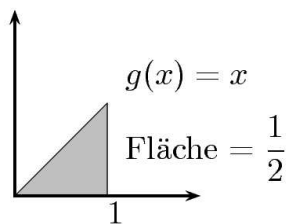
(b) Paraboloid



$f(x) = \sqrt{x}$ im Intervall $[0, 1]$:

$$\begin{aligned} V &= \pi \int_0^1 (f(x))^2 \, dx \\ &= \pi \int_0^1 x \, dx \\ &= \frac{\pi}{2} \end{aligned}$$

denn $\int_0^1 x \, dx = \frac{1}{2}$.



Kapitel 29

Das unbestimmte Integral und die Stammfunktion

29.1 Motivation

- Gibt es eine Operation, die die Differentiation „rückgängig“ macht?
- Kann man zu einer gegebenen Funktion f eine Funktion F finden mit $F'(x) = f(x)$?

29.2 Definition: (Stammfunktion)

Gegeben seien Funktionen $F : [a, b] \rightarrow \mathbb{R}$. Ist F differenzierbar auf $[a, b]$ und gilt $F'(x) = f(x)$ für alle $x \in [a, b]$, so heißt $F(x)$ Stammfunktion.

Bemerkung: Ist $F(x)$ eine Stammfunktion von $f(x)$, so ist auch $F(x) + C$ mit einer beliebigen Konstante C eine Stammfunktion von $f(x)$.

Kommen wir nun zum Gipfel der Analysis.

29.3 Satz: (Hauptsatz der Differential- und Integralrechnung)

Sei $f : [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion. Dann gilt:

- (a) $F(x) := \int_a^x f(t) \, dt$ ist eine Stammfunktion von $f(x)$
(„Integration als Umkehrung der Differentiation“).
- (b) Ist $F(x)$ Stammfunktion von $f(x)$, so gilt

$$\int_a^b f(t) \, dt = F(b) - F(a) =: F(x)|_a^b$$

(„Stammfunktion als Mittel zum Lösen bestimmter Integrale“).

Beweis:

(a) Sei $h \neq 0$ ($h < 0$ ebenfalls zugelassen) so, dass $x, x + h \in [a, b]$.

$$\begin{aligned} \left| \frac{1}{h} (F(x+h) - F(x) - f(x)h) \right| &= \left| \frac{1}{h} \left(\int_a^{x+h} f(t) \, dt - \int_a^x f(t) \, dt \right) - \frac{1}{h} \int_x^{x+h} \underbrace{f(x)}_{\text{unabhängig von } t} \, dt \right| \\ &= \frac{1}{|h|} \left| \int_x^{x+h} (f(t) - f(x)) \, dt \right| \\ &\leq \sup \{ |f(t) - f(x)| \mid t \text{ zwischen } x \text{ und } x+h \} \\ &\rightarrow 0 \text{ für } h \rightarrow 0 \text{ da } f \text{ stetig.} \end{aligned}$$

Somit ist $F'(x) = f(x)$.

(b) Mit (a) gilt:

$$F(x) = \int_a^x f(t) \, dt + C$$

Damit folgt:

$$\begin{aligned} F(b) &= \int_a^b f(t) \, dt + C \\ F(a) &= \int_a^a f(t) \, dt + C = C \\ \Rightarrow F(b) - F(a) &= \int_a^b f(t) \, dt. \end{aligned}$$

□

29.4 Definition: (Unbestimmte Integral)

Die Stammfunktion $F(x)$ einer Funktion $f(x)$ wird auch als „das“ unbestimmte Integral von $f(x)$ bezeichnet und $\int f(x) \, dx$ (also ohne Integrationsgrenzen) geschrieben. Es ist jedoch nur bis auf eine additive Konstante C bestimmt.

29.5 Beispiele von unbestimmten Integralen

$$\begin{aligned}\int x^n dx &= \frac{1}{n+1} x^{n+1} + C & (n \neq -1) \\ \int \frac{1}{x} dx &= \ln|x| + C & (x \neq 0) \\ \int \sin x dx &= -\cos x + C \\ \int \cos x dx &= \sin x + C \\ \int \tan x dx &= -\ln|\cos x| + C & (\cos x \neq 0) \\ \int \frac{1}{\cos^2 x} dx &= \tan x + C & (x \neq (2k+1) \cdot \frac{\pi}{2}, k \in \mathbb{Z}) \\ \int e^{ax} dx &= \frac{1}{a} e^{ax} + C & (a \neq 0) \\ \int \ln x dx &= x(\ln x - 1) + C & (x > 0) \\ \int \frac{1}{\sqrt{1-x^2}} dx &= \arcsin x + C & (|x| < 1) \\ \int \frac{1}{\sqrt{x^2-1}} dx &= \ln|x + \sqrt{x^2-1}| + C & (|x| > 1) \\ \int \frac{1}{\sqrt{x^2+1}} dx &= \ln|x + \sqrt{x^2+1}| + C \\ \int \frac{1}{1+x^2} dx &= \arctan x + C \\ \int \frac{1}{1-x^2} dx &= \frac{1}{2} \ln \left| \frac{1+x}{1-x} \right| + C & (|x| \neq 1)\end{aligned}$$

weitere Beispiele: Formelsammlungen, z.B. Bronstein: Taschenbuch der Mathematik.

Unbekannte Integrale berechnet man durch (z.T. trickreiche) Umformungen in bekannte Integrale. Hierzu gibt es zwei Grundtechniken:

29.6 Satz: (Integrationsregeln)

(a) Partielle Integration:

Sind $u, v : [a, b] \rightarrow \mathbb{R}$ stetig differenzierbar, so gilt:

$$\int u(x)v'(x) dx = u(x)v(x) - \int u'(x)v(x) dx$$

und für bestimmte Integrale:

$$\int_a^b u(x)v'(x) dx = u(x)v(x)|_a^b - \int_a^b u'(x)v(x) dx.$$

(b) Substitutionsregel:

Zur Berechnung von $\int_c^d f(x) \, dx$ setze man $x = g(t)$ mit g stetig differenzierbar, $g(a) = c$, $g(b) = d$ und formal $dx = g'(t)dt$:

$$\int_{x=c}^{x=d} f(x) \, dx = \int_{t=a}^{t=b} f(g(t))g'(t) \, dt.$$

Beweis:

(a) folgt aus der Produktregel der Differentiation:

$$\begin{aligned} (uv)' &= u'v + uv' \\ \Rightarrow \underbrace{\int_a^b (uv)' \, dx}_{uv|_a^b} &= \int_a^b u'v \, dx + \int_a^b v'u \, dx. \end{aligned}$$

(b) folgt aus der Kettenregel

$$\frac{d}{dt}(F(g(t))) = F'(g(t)) \cdot g'(t) = f(g(t)) \cdot g'(t).$$

Integration über t im Intervall $[a, b]$ ergibt:

$$\begin{aligned} \int_a^b f(g(t))g'(t) \, dt &= F(g(t))|_a^b = F(g(b)) - F(g(a)) \\ &= \int_{g(a)}^{g(b)} f(x) \, dx. \end{aligned}$$

□

29.7 Beispiele zur partiellen Integration

(a)

$$\begin{aligned} \int \underbrace{x}_u \underbrace{e^x}_{v'} \, dx &= \underbrace{x}_u \underbrace{e^x}_v - \int \underbrace{1}_{u'} \cdot \underbrace{e^x}_v \, dx \\ &= xe^x - e^x + C \end{aligned}$$

(b)

$$\begin{aligned} \int \ln x \, dx &= \int \underbrace{1}_{v'} \cdot \underbrace{\ln x}_u \, dx \\ &= x \ln x - \int x \cdot \frac{1}{x} \, dx = x \ln x - x + C \end{aligned}$$

(c)

$$\begin{aligned}
 \int \sin^2 x \, dx &= \int \underbrace{\sin x}_u \cdot \underbrace{\sin x}_{v'} \, dx \\
 &= \sin x (-\cos x) + \int \cos^2 x \, dx \\
 &= -\sin x \cos x + \int (1 - \sin^2 x) \, dx \quad | + \int \sin^2 x \, dx \\
 \Rightarrow 2 \int \sin^2 x \, dx &= -\sin x \cos x + \int dx \\
 \int \sin^2 x \, dx &= -\frac{1}{2} \sin x \cos x + \frac{x}{2} + C
 \end{aligned}$$

29.8 Beispiele zur Substitutionsregel

(a) $\int e^{\sqrt{x}} \, dx$

Subst. $z = \sqrt{x}, \frac{dz}{dx} = \frac{1}{2\sqrt{x}} = \frac{1}{2z}$

$dx = 2z \, dz$

$$\begin{aligned}
 \int e^{\sqrt{x}} \, dx &= \int e^z \cdot 2z \, dz \\
 &= 2 \cdot \int z e^z \, dz \\
 &\stackrel{29.7(a)}{=} 2(z e^z - e^z) + C \\
 &= 2(\sqrt{x} e^{\sqrt{x}} - e^{\sqrt{x}}) + C
 \end{aligned}$$

(b) $\int_{-1}^1 \sqrt{1-x^2} \, dx$

Subst. $x = \cos z, \quad \frac{dx}{dz} = -\sin z, \, dx = -\sin z \, dz$

$$\begin{aligned}
 \int_{-1}^1 \sqrt{1-x^2} \, dx &= \int_{z=\pi}^{z=0} \sqrt{1-\cos^2 z} \cdot (-\sin z) \, dz \\
 &= -\int_{\pi}^0 \sin^2 z \, dz = \int_{\pi}^0 \sin^2 z \, dz \\
 &\stackrel{29.7(c)}{=} -\frac{1}{2} \sin z \cos z + \frac{z}{2} \Big|_{\pi}^0 \\
 &= -0 + \frac{\pi}{2} - (-0 + 0) = \frac{\pi}{2}.
 \end{aligned}$$

29.9 Bemerkung:

Nicht alle Integrale lassen sich analytisch lösen, obwohl sie z.T. einfach aussehen.
Beispiel „Fehlerfunktion“ $\operatorname{erf}(x) := \frac{2}{\sqrt{\pi}} \int_0^x e^{-t^2} dt$ ist in Formelsammlungen tabelliert bzw. numerisch approximierbar.

29.10 Integration rationaler Funktionen

Rückführung auf bekannte Integrale mittels Partialbruchzerlegung.

Beispiel: $\int \frac{1}{x^2 - 5x + 6} dx$

Wegen $x^2 - 5x + 6 = (x - 2)(x - 3)$ machen wir den Ansatz

$$\begin{aligned} \frac{1}{x^2 - 5x + 6} &= \frac{A}{x - 2} + \frac{B}{x - 3} \\ &= \frac{A(x - 3) + B(x - 2)}{(x - 2)(x - 3)} \\ &= \frac{(A + B)x + (-3A - 2B)}{x^2 - 5x + 6} \end{aligned}$$

Koeffizientenvergleich:

$$\left. \begin{array}{lcl} (1) & A + B & = 0 \\ (2) & -3A - 2B & = 1 \end{array} \right\} \text{Lösung: } A = -1, B = 1.$$

Addition des 3-fachen von (1) zu (2) ergibt $B = 1$; in (1): $A = -1$.

$$\begin{aligned} \Rightarrow \frac{1}{x^2 - 5x + 6} &= \frac{-1}{x - 2} + \frac{1}{x - 3} \\ \int \frac{1}{x^2 - 5x + 6} dx &= -\int \frac{1}{x - 2} dx + \int \frac{1}{x - 3} dx \\ &= -\ln|x - 2| + \ln|x - 3| + C \\ &= \ln \left| \frac{x - 3}{x - 2} \right| + C \quad (x \neq 2, 3) \end{aligned}$$

Kapitel 30

Uneigentliche Integrale

30.1 Motivation

Manchmal muss man Integrale berechnen

- bei denen über ein unendliches Intervall integriert wird.

Beispiel: $\int_a^\infty f(x) \, dx$, $\int_{-\infty}^b f(x) \, dx$, $\int_{-\infty}^\infty f(x) \, dx$

- bei denen unbeschränkte Funktionen vorliegen.

Beispiel: $\int_0^1 \frac{1}{\sqrt{x}} \, dx$

Solche Integrale heißen uneigentliche Integrale.

Unter welchen Bedingungen ist deren Berechnung möglich?

30.2 Definition: (Unendliche Integrationsgrenzen)

Sei $f : [a, \infty) \rightarrow \mathbb{R}$ über jedem Intervall $[a, R]$ mit $a < R < \infty$ integrierbar.

Falls $\lim_{R \rightarrow \infty} \int_a^R f(x) \, dx$ existiert, heißt $\int_a^\infty f(x) \, dx$ konvergent und man setzt

$$\int_a^\infty f(x) \, dx := \lim_{R \rightarrow \infty} \int_a^R f(x) \, dx.$$

Analog definiert man $\int_{-\infty}^b f(x) \, dx$ für $f : (-\infty, b] \rightarrow \mathbb{R}$.

Ferner setzt man

$$\int_{-\infty}^\infty f(x) \, dx := \int_{-\infty}^a f(x) \, dx + \int_a^\infty f(x) \, dx \text{ für } f : \mathbb{R} \rightarrow \mathbb{R}.$$

30.3 Beispiele

(a) Konvergiert $\int_1^\infty \frac{1}{x^s} dx$ für $s > 1$?

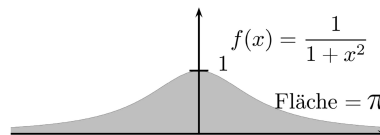
$$\begin{aligned} \int_1^\infty \frac{1}{x^s} dx &= \left. \frac{1}{1-s} x^{1-s} \right|_1^R \\ &= \frac{1}{1-s} \left(\underbrace{R^{1-s}}_{\rightarrow 0} - 1^{1-s} \right) \rightarrow \frac{1}{s-1} \text{ für } R \rightarrow \infty. \end{aligned}$$

Also ist $\int_1^\infty \frac{1}{x^s} dx$ konvergent für $s > 1$.

Bemerkung: Für $s \leq 1$ ist $\int_1^\infty \frac{1}{x^s} dx$ divergent.

(b)

$$\begin{aligned} \int_{-\infty}^\infty \frac{1}{1+x^2} dx &= \lim_{R \rightarrow \infty} \int_{-R}^0 \frac{1}{1+x^2} dx + \lim_{R \rightarrow \infty} \int_0^R \frac{1}{1+x^2} dx \\ &= \lim_{R \rightarrow \infty} \arctan x \Big|_{-R}^0 + \lim_{R \rightarrow \infty} \arctan x \Big|_0^R \\ &= 0 - \left(-\frac{\pi}{2}\right) + \frac{\pi}{2} - 0 = \pi \end{aligned}$$



Integration unbeschränkter Funktionen:

Wie geht man vor, wenn der Integrand an einer der Integrationsgrenzen nicht definiert ist?

Beispiel: $\int_0^1 \frac{1}{\sqrt{x}} dx$

30.4 Definition: (Konvergenz)

Sei $f : (a, b] \rightarrow \mathbb{R}$ über jedem Teilintervall $[a+\epsilon, b]$ mit $0 < \epsilon < b-a$ integrierbar und in a nicht definiert.

Falls $\lim_{\epsilon \rightarrow 0} \int_{a+\epsilon}^b f(x) dx$ existiert, heißt $\int_a^b f(x) dx$ konvergent, und man setzt

$$\int_a^b f(x) dx := \lim_{\epsilon \rightarrow 0} \int_{a+\epsilon}^b f(x) dx.$$

Eine analoge Definition gilt, falls $f : [a, b) \rightarrow \mathbb{R}$ in b nicht definiert ist:

$$\int_a^b f(x) \, dx := \lim_{\epsilon \rightarrow 0} \int_a^{b-\epsilon} f(x) \, dx.$$

Falls $f : (a, b) \rightarrow \mathbb{R}$ in a und b nicht definiert ist, setzt man

$$\int_a^b f(x) \, dx := \int_a^c f(x) \, dx + \int_c^b f(x) \, dx \quad \text{mit } a < c < b.$$

30.5 Beispiel:

Konvergiert $\int_0^1 \frac{1}{x^s} \, dx$ für $0 < s < 1$?

$$\begin{aligned} \int_{\epsilon}^1 \frac{1}{x^s} \, dx &= \left. \frac{1}{1-s} x^{1-s} \right|_{\epsilon}^1 \\ &= \frac{1}{1-s} (1 - \epsilon^{1-s}) \rightarrow \frac{1}{1-s} \text{ für } \epsilon \rightarrow 0 \\ &\Rightarrow \text{Konvergenz} \end{aligned}$$

Bemerkung: $\int_0^1 \frac{1}{x^s} \, dx$ konvergiert nicht für $s \geq 1$:

$$\begin{aligned} s = 1 : \quad \int_{\epsilon}^1 \frac{1}{x} \, dx &= \left. \ln|x| \right|_{\epsilon}^1 \\ &= \underbrace{\ln 1}_0 - \ln \epsilon \rightarrow \infty \text{ für } \epsilon \rightarrow 0. \\ s > 1 : \quad \int_{\epsilon}^1 \frac{1}{x^s} \, dx &= \left. \frac{1}{1-s} x^{1-s} \right|_{\epsilon}^1 \\ &= \frac{1}{s-1} \left(\frac{1}{\epsilon^{s-1}} - 1 \right) \rightarrow \infty \text{ für } \epsilon \rightarrow 0. \end{aligned}$$

30.6 Bemerkung:

Falls Polstellen im Innern des Integrationsbereich vorliegen, spaltet man das Integral so in Teilintegrale auf, dass die Polstellen an den Grenzen der Teilintegrale liegen.

Für uneigentlich Integrale kann man ähnliche Konvergenzkriterien wie für Reihen zeigen (vgl. ¶ 16 auf Seite 123).

30.7 Satz: (Konvergenzkriterien für uneigentliche Integrale)

Sei $f : [a, \infty) \rightarrow \mathbb{R}$ integrierbar auf jedem endlichen Intervall $[a, b]$. Dann gilt:

(a) **Cauchy-Kriterium:**

$$\int_a^\infty f(x) \, dx \text{ existiert} \Leftrightarrow \forall \epsilon > 0 \, \exists c > a : \left| \int_{z_1}^{z_2} f(x) \, dx \right| < \epsilon \quad \forall z_1, z_2 > c.$$

(b) **Absolute Konvergenz:**

Ist $\int_a^\infty |f(x)| \, dx$ (d.h. $\int_a^\infty f(x) \, dx$ konvergiert absolut), so konvergiert auch $\int_a^\infty f(x) \, dx$.

(c) **Majorantenkriterium:** Ist $|f(x)| \leq g(x)$ für alle $x \in [a, \infty)$ und konvergiert $\int_a^\infty g(x) \, dx$, so konvergiert $\int_a^\infty f(x) \, dx$ absolut. Gilt umgekehrt $0 \leq g(x) \leq f(x)$ und divergiert $\int_a^\infty g(x) \, dx$, so divergiert auch $\int_a^\infty f(x) \, dx$.

Bemerkung: Entsprechende Aussagen gelten auch beim nach unten unbeschränkten Integrationsintervall $(-\infty, b]$, sowie an Singularitäten an den Grenzen eines beschränkten Intervalls $[a, b]$.

30.8 Beispiel: Das Dirichlet-Integral $\int_0^\infty \frac{\sin x}{x} \, dx$

Cauchy-Kriterium ergibt für $0 < z_1 < z_2$:

$$\begin{aligned} \int_{z_1}^{z_2} \frac{\sin x}{x} \, dx & \stackrel{\text{part. Int.}}{=} -\frac{\cos x}{x} \Big|_{z_1}^{z_2} - \int_{z_1}^{z_2} \frac{\cos x}{x^2} \, dx \\ \Rightarrow \left| \int_{z_1}^{z_2} \frac{\sin x}{x} \, dx \right| & \leq \frac{1}{z_1} + \frac{1}{z_2} + \int_{z_1}^{z_2} \frac{1}{x^2} \, dx \rightarrow 0 \text{ für } z_1 \rightarrow \infty \end{aligned}$$

(da $\int_1^\infty \frac{1}{x^2} \, dx$ nach 30.3 (a) konvergiert).

Das Dirichlet-Integral tritt z.B. in der Signalverarbeitung auf.

30.9 Beispiel: Die Gammafunktion (Euler, 1707-1783)

$$\Gamma(x) := \int_0^\infty e^{-t} t^{x-1} \, dt \quad \text{für } x > 0$$

Konvergiert dieses Integral?

Probleme:

- (a) Für $0 < x < 1$ divergiert der Integrand in $t = 0$.
- (b) Obere Integrationsgrenze ist unendlich.

Wir spalten daher auf

$$\Gamma(x) := \int_0^1 e^{-t} t^{x-1} dt + \int_1^\infty e^{-t} t^{x-1} dt$$

und behandeln beide Probleme getrennt:

- (a) Betrachte $\int_0^1 e^{-t} t^{x-1} dt$ für $0 < x < 1$.

Wegen $0 \leq e^{-t} t^{x-1} \leq t^{x-1}$ hat $\int_0^1 e^{-t} t^{x-1} dt$ die nach 30.5 konvergente

Majorante $\int_0^1 \frac{1}{t^{1-x}} dt$.

- (b) Betrachte $\int_1^\infty e^{-t} t^{x-1} dt$

Es existiert $c > 0$ mit $0 \leq e^{-t} t^{x-1} \leq \frac{c}{t^2}$, denn

$$ce^t \geq t^{x+1} \quad (e - \text{Funktion wächst schneller als jede Potenz}).$$

Somit ist $\frac{c}{t^2}$ konvergente Majorante (vgl. 30.3 (a)).

Also existiert $\Gamma(x)$ für all $x > 0$.

Zwei wichtige Eigenschaften der Gammafunktion:

(i)

$$\begin{aligned} \Gamma(1) &= \int_0^\infty e^{-t} dt \\ &= -e^{-t} \Big|_0^\infty = -0 + 1 = 1 \end{aligned}$$

$$\boxed{\Gamma(1) = 1}$$

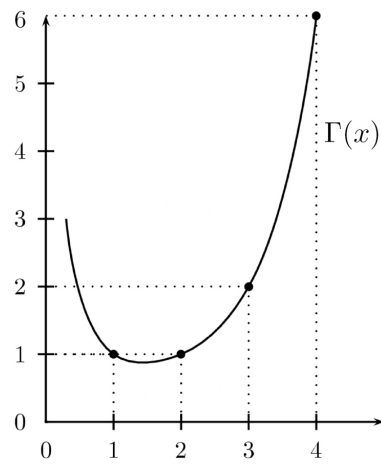
(ii)

$$\begin{aligned} \Gamma(x) &= \int_0^\infty \underbrace{e^{-t}}_u \underbrace{t^{x-1}}_{v'} dt \\ &= \underbrace{e^{-t} \frac{1}{x} t^x \Big|_{t=0}^{t=\infty}}_0 + \frac{1}{x} \int_0^\infty e^{-t} t^x dt \\ &= \frac{1}{x} \Gamma(x+1) \end{aligned}$$

d.h. $\Gamma(x+1) = x \cdot \Gamma(x)$

Wegen (i) und (ii) interpoliert die Gammafunktion die Fakultät.

$$\Gamma(n+1) = n!$$



Kapitel 31

Numerische Verfahren zur Integration

31.1 Motivation

Die wenigsten Integrale $\int_a^b f(x) \, dx$ lassen sich mit Hilfe von Stammfunktionen analytisch berechnen. Man ist deswegen auf numerische Verfahren, sogenannte Quadraturverfahren angewiesen.

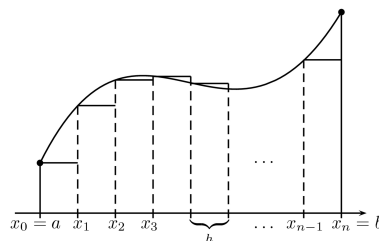
Gesucht:

$I[f] := \int_a^b f(x) \, dx$ für eine Funktion $f(x)$, die auf $[a, b]$ hinreichend oft differenzierbar ist.

Quadraturformeln führen dieses Integral in eine Summe über:

$$\begin{aligned} I[f] &\approx \sum_{i=0}^n w_i f(x_i) : && \text{Quadraturformel} \\ x_i &\in [a, b], \quad i = 0, 1, \dots, n : && \text{Knoten} \\ w_i, i &= 0, \dots, n : && \text{Gewichte} \end{aligned}$$

Praktisch brauchbar sind nur Quadraturformeln mit nichtnegativen Gewichten. Sie garantieren Nichtnegativität der Integralapproximation für nichtnegative Integranden, sowie gutartiges Verhalten gegenüber Rundungsfehlern.



31.2 Einfachstes Beispiel: Rechteckregel

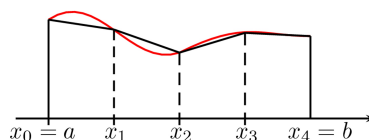
$$\begin{aligned} x_i &= a + i \cdot \frac{b-a}{n}, \quad i = 0, \dots, n, \text{ Knoten} \\ h &= \frac{b-a}{n}, \quad \text{Gitterweite} \\ R_h[f] &:= \sum_{i=0}^{n-1} h \cdot f(x_i) \end{aligned}$$

Im Allgemeinen ist die Rechteckregel nicht sehr genau, solange man nicht sehr viele Knoten verwendet.

Grund: $f(x)$ wird im Intervall $[x_i, x_{i+1}]$ durch die Konstante $f(x_i)$ approximiert.

31.3 Zweiteinfachtes Beispiel: Trapezregel

Idee: Approximiere $f(x)$ in $[x_i, x_{i+1}]$ durch Gerade durch Punkte $(x_i, f(x_i))$ und $(x_{i+1}, f(x_{i+1}))$.



Fläche aller Trapeze:

$$\begin{aligned} T_h[f] &= \sum_{i=0}^{n-1} h \cdot \frac{f(x_i) + f(x_{i+1})}{2} \\ &= \frac{h}{2} f(x_0) + h \cdot (f(x_1) + f(x_2) + \dots + f(x_{n-1})) + \frac{h}{2} f(x_n). \end{aligned}$$

31.4 Bemerkung:

Genauere Quadraturformeln erhält man, indem $f(x)$ stückweise durch Polynome höheren Grades ersetzt wird ($\rightarrow$ **Newton-Cotes-Verfahren**).

Für zu große Grade treten allerdings negative Gewichte auf und die Formeln werden unbrauchbar.

Gewichte der Newton-Cotes-Formeln

Polynomgrad	Gewichte pro Intervall					Name
0	1					Rechteckregel
1	$\frac{1}{2}$		$\frac{1}{2}$			Trapezregel
2	$\frac{1}{6}$		$\frac{4}{6}$	$\frac{1}{6}$		Simpson-Regel
3	$\frac{1}{8}$	$\frac{3}{8}$	$\frac{3}{8}$	$\frac{1}{8}$		$\frac{3}{8}$ -Regel
4	$\frac{7}{90}$	$\frac{32}{90}$	$\frac{12}{90}$	$\frac{32}{90}$	$\frac{7}{90}$	Milue-Regel

Die Newton-Cotes-Formel integriert Polynome $\leq n$ exakt.

Gibt es bessere Alternativen zur Erhöhung der Genauigkeit?

Anwort: **Romberg-Quadratur**

Idee: Fehlerterme in der Trapezregel sukzessive wegextrapolieren.
Es gilt (Beweis aufwändig):

31.5 Satz: (Euler-Maclaurinsche Summenformel)

Für $f \in C^\infty([a, b])$ besitzt die Trapezregel die asymptotische Fehlerentwicklung

$$T_h[f] = \int_a^b f(x) \, dx + c_1 \cdot h^2 + c_2 \cdot h^4 + c_3 \cdot h^6 + \dots$$

mit Konstanten c_k , die von $f^{(2k-1)}(a)$ und $f^{(2k-1)}(b)$ abhängen.

Konsequenz:

$$\begin{aligned} T_h[f] &= \int_a^b f(x) \, dx + c_1 \cdot h^2 + c_2 \cdot h^4 + \mathcal{O}(h^6) \\ T_{\frac{h}{2}}[f] &= \int_a^b f(x) \, dx + c_1 \cdot \frac{h^2}{4} + c_2 \cdot \frac{h^4}{16} + \mathcal{O}(h^6) \\ \Rightarrow \frac{4 \cdot T_{\frac{h}{2}}[f] - T_h[f]}{3} &= \int_a^b f(x) \, dx - \frac{1}{4} c_2 \cdot h^4 + \mathcal{O}(h^6) \end{aligned}$$

Quadratischer Fehlerterm in h wurde wegextrapoliert. Ebenso kann man mit $T_h[f], T_{\frac{h}{2}}[f], T_{\frac{h}{4}}[f]$ die h^4 -Terme eliminieren und erhält ein Verfahren 6. Ordnung!

31.6 Satz: (Romberg-Quadratur)

Es sei $f \in C^{2m+2}([a, b])$ und $T_h[f]$ die Trapezapproximation an $\int_a^b f(x) \, dx$ mit Schrittweite h .

Dann erhält man mit der Rekursionsformel

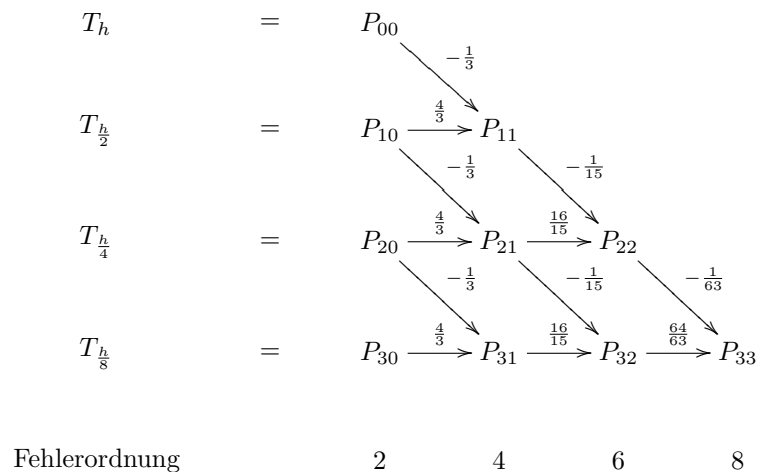
$$\begin{aligned} P_{k0} &:= T_{\frac{h}{2^k}}[f], & (k = 0, 1, \dots, n) \\ P_{kj} &:= \frac{4^j \cdot P_{k,j-1} - P_{k-1,j-1}}{4^j - 1}, & \begin{matrix} k = 1, \dots, m \\ j = 1, \dots, m \end{matrix} \end{aligned}$$

in $P_{m,m}$ eine Quadraturformel der Fehlerordnung $2m + 2$.

Beweis:

Übungsaufgabe (vollst. Ind.)

Veranschaulichung durch Schema:



Verfahrensmerkmale:

- schrittweise Intervallhalbierung
- Werte aus vorangegangenen Schritten werden weiter verwendet.
- Abbruch der Verfeinerung, Approximationen unterscheiden sich hinreichend dicht voneinander.

31.7 Beispiel:

$$\int_0^1 \frac{\sin t}{t} \, dt$$

h	$T_h[f]$	Romberg-Quadratur	#Funktionsauswertung
1	0.9207354924	0.9207354924	2
$\frac{1}{2}$	0.9397932848	0.9461458823	3
$\frac{1}{4}$	0.9445135217	0.9460830041	5
$\frac{1}{8}$	0.9456908636	0.9460830704	9
$\frac{1}{16}$	0.9459850299	0.9460830704	17
$\vdots$	$\vdots$	$\vdots$	
$\frac{1}{32768}$	0.9460830703		32769

Fazit: Die Romberg-Quadratur hat den Aufwand gegenüber der Trapezregel um den Faktor 2000 reduziert, bei vergleichbarer Genauigkeit.

Kapitel 32

Kurven und Bogenlänge

32.1 Motivation

Der Begriff der Kurve in der Ebene oder im Raum spielt in den Naturwissenschaften, insbesondere der Physik, Technik (Robotik) und der Informatik (Computergraphik) eine tragende Rolle.

Neben der Berechnung von Flächen und Volumina ist die Definition / Ermittlung der Länge von Kurven eine weitere wichtige Anwendung der Integralrechnung.

32.2 Definition: (Parametrisierte Kurve)

Eine (parametrisierte) Kurve im $\mathbb{R}^n$ ist eine Abbildung

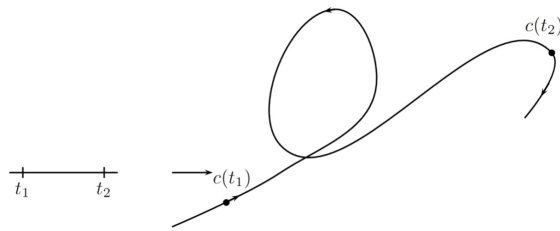
$$\begin{aligned} c : [a, b] &\rightarrow \mathbb{R}^n \\ t &\mapsto (x_1(t), \dots, x_n(t))^T := \begin{pmatrix} x_1(t) \\ \vdots \\ x_n(t) \end{pmatrix}, \end{aligned}$$

deren Komponentenfunktionen $x_1, \dots, x_n : [a, b] \rightarrow \mathbb{R}$ stetig sind. c heißt differenzierbar (stetig differenzierbar, C^1 -Kurve), wenn alle x_i differenzierbar (stetig differenzierbar) sind. Man definiert

$$\dot{c}(t) := \left(\frac{dx_1}{dt}(t), \dots, \frac{dx_n}{dt}(t) \right)^T.$$

Das Bild $c([a, b])$ heißt Spur von c .

„Bewegung eines Punktes im Raum, $c(t)$ als Ort des Punktes zur Zeit t “.



32.3 Beispiele

- (a) Ellipse mit Hauptachsen a und b :

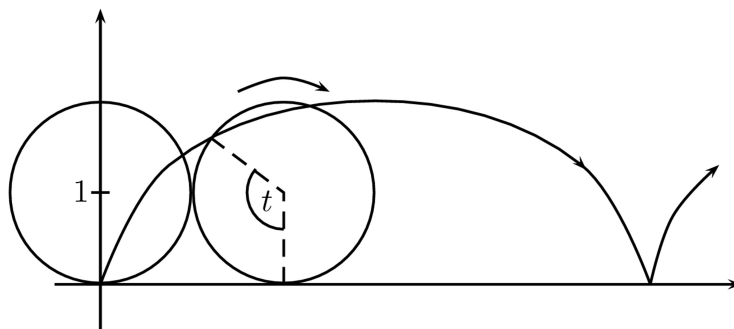
$$x(t) = a \cdot \cos t$$

$$y(t) = b \cdot \sin t$$

Aus $\cos^2 t + \sin^2 t = 1$ folgt die Spurgleichung

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1.$$

- (b) Zykloide



Kurve eines Randpunktes eines Kreises, der auf einer Geraden abrollt.

$$x(t) = t - \sin(t)$$

$$y(t) = 1 - \cos(t)$$

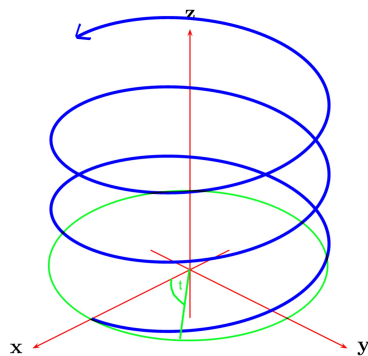
Überlagerung des Mittelpunktsbewegung $\begin{pmatrix} x(t) \\ y(t) \end{pmatrix} = \begin{pmatrix} t \\ 1 \end{pmatrix}$

und einer Kreisbewegung $\begin{pmatrix} x(t) \\ y(t) \end{pmatrix} = \begin{pmatrix} -\sin(t) \\ -\cos(t) \end{pmatrix}$ im Uhrzeigersinn.

- (c) Schraubenlinie

$$c(t) = (r \cdot \cos t, r \cdot \sin t, h \cdot t)^T, \quad t \in \mathbb{R}.$$

Überlagerung einer Kreisbewegung mit Radius r in der $x-y$ -Ebene und einer linearen Bewegung in z -Richtung.



32.4 Definition: (Parameterwechsel, Umparametrisierung)

Ist $c : [a, b] \rightarrow \mathbb{R}^n$ eine Kurve und $h : [\alpha, \beta] \rightarrow [a, b]$ eine stetige, bijektive und monoton wachsende Abbildung, so hat die „neue“ Kurve $\tilde{c}(\tau) := c(h(\tau))$, $\tau \in [\alpha, \beta]$ die selbe Spur und den selben Durchlaufsinne wie c . Man nennt $t = h(\tau)$ einen Parameterwechsel (oder Umparametrisierung).

Kurven, die durch einen Parameterwechsel auseinander hervorgehen, werden als gleich angesehen.

Ist c stetig differenzierbar, so werden nur stetig differenzierbare Funktionen $h : [\alpha, \beta] \rightarrow [a, b]$ mit $h'(\tau) > 0$ als Parameterwechsel zugelassen (C^1 -Parameterwechsel).

32.5 Bemerkung:

Verschiedene Kurvenparametrisierungen können zur selben Spur führen:

$$\begin{aligned} c_1(t) &:= (\cos t, \sin t), & t \in [0, 2\pi] \\ \text{und} \\ c_2(t) &:= (\cos t, -\sin t), & t \in [0, 2\pi] \end{aligned}$$

haben den Einheitskreis als Spur, unterscheiden sich aber im Durchlaufsinne.

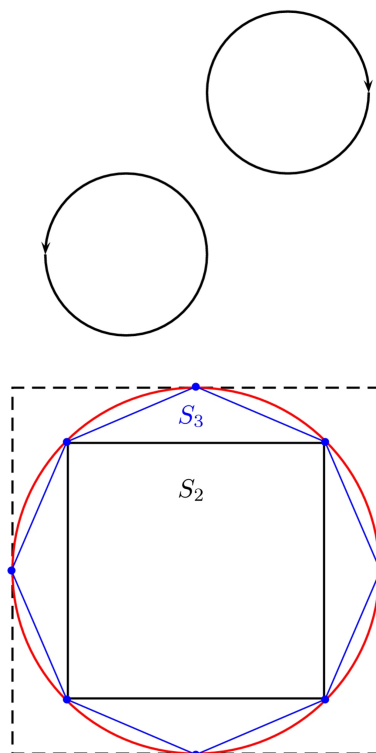
Bogenlänge einer Kurve

wichtige Anwendung der Integralrechnung auf Kurven.

Motivierendes Beispiel: (Kreisumfang)

Approximiere Kreis durch einbeschriebene regelmäßige 2^n -Ecke ($n \geq 2$). Ihre Umfänge S_n wachsen monoton in n und sind z.B. durch den Umfang eines umbeschriebenen Quadrats nach oben beschränkt. Somit existiert $\sup_n(S_n)$.

Es definiert den Kreisumfang.
Allgemein:



Approximiere die Kurve $c(t)$ durch einen Polygonzug. Für eine Zerlegung

$$Z = \{a = t_0 < t_1 < \dots < t_m = b\}$$

von $[a, b]$ ist seine Länge gegeben durch

$$L(Z) = \sum_{i=0}^{m-1} |c(t_{i+1}) - c(t_i)|.$$

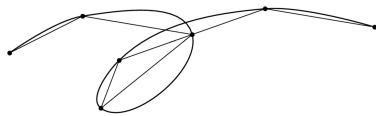
Dabei bezeichnet $|(x_1, \dots, x_n)^T| := \sqrt{\sum_{i=1}^n x_i^2}$ die euklidische Norm eines Vektors $(x_1, \dots, x_n)^T$.

32.6 Definition: (Länge, Rektifizierbarkeit einer Kurve)

Es bezeichne $\mathfrak{Z}[a, b]$ die Menge aller Zerlegungen von $\mathfrak{Z}[a, b]$. Ist die Menge $\{L(Z) \mid Z \in \mathfrak{Z}[a, b]\}$ nach oben beschränkt, so heißt die Kurve c rektifizierbar, und

$$L(c) := \sup\{L(Z) \mid Z \in \mathfrak{Z}[a, b]\}$$

heißt die Länge der Kurve c .



32.7 Satz:

Jede C^1 -Kurve $c : [a, b] \rightarrow \mathbb{R}^n$ ist rektifizierbar, und es gilt

$$L(c) := \int_a^b |\dot{c}(t)| \, dt.$$

Beweis:

Für eine Zerlegung Z gilt:

$$\begin{aligned} L(Z) &= \sum_{i=0}^{m-1} |c(t_{i+1}) - c(t_i)| \\ &= \sum_{i=0}^{m-1} \sqrt{\sum_{k=1}^n (x_k(t_{i+1}) - x_k(t_i))^2} && \text{Def. der euklid. Norm} \\ &= \sum_{i=0}^{m-1} \sqrt{\sum_{k=1}^n (\dot{x}_k(\tau_{k_i}))^2 (t_{i+1} - t_i)} && \text{Mittelwertsatz} \end{aligned}$$

mit $\tau_{k_i} \in [t_i, t_{i+1}]$.

Für die zu Z gehörende Rechtecksumme $R(Z)$ von $\int_a^b |\dot{c}(t)| \, dt$ gilt:

$$R(Z) = \sum_{i=0}^{m-1} \sqrt{\sum_{k=1}^n (\dot{x}_k(t_i))^2} (t_{i+1} - t_i)$$

Zur Abschätzung von $|L(Z) - R(Z)|$ verwenden wir die gleichmäßige Stetigkeit der $\dot{x}_k(t)$ auf dem Kompaktum $[a, b]$:

$$\forall \varepsilon > 0 \, \exists \delta > 0 : |\tilde{t} - t| < \delta \Rightarrow |\dot{x}_k(\tilde{t}) - \dot{x}_k(t)| < \varepsilon, \quad k = 1, \dots, n.$$

Gilt nun für gegebenes $\epsilon > 0$, dass die Feinheit $\|Z\|$ der Zerlegung $\|Z\| < \delta$

erfüllt, so folgt damit

$$\begin{aligned}
 |L(Z) - R(Z)| &= \left| \sum_{i=0}^{m-1} \left(\sqrt{\sum_{k=1}^n (\dot{x}_k(\tau_{k_i}))^2} - \sqrt{\sum_{k=1}^n (\dot{x}_k(t_i))^2} \right) \cdot (t_{i+1} - t_i) \right| \\
 &\stackrel{\Delta-Ungl.}{\leq} \sum_{i=0}^{m-1} \left| \sqrt{\sum_{k=1}^n (\dot{x}_k(\tau_{k_i}))^2} - \sqrt{\sum_{k=1}^n (\dot{x}_k(t_i))^2} \right| \cdot (t_{i+1} - t_i) \\
 &\stackrel{(*)}{\leq} \sum_{i=0}^{m-1} \sqrt{\sum_{k=1}^n (\dot{x}_k(\tau_{k_i}) - \dot{x}_k(t_i))^2} \cdot (t_{i+1} - t_i) \\
 &\stackrel{(*)}{=} \text{denn man zeigt für die eukl. Norm: } ||a| - |b|| \leq |a - b| \\
 &\leq \sum_{i=0}^{m-1} \underbrace{\sqrt{\sum_{k=1}^n \varepsilon^2}}_{=\sqrt{n} \cdot \varepsilon} \cdot (t_{i+1} - t_i) \\
 &= \sqrt{n} \cdot \varepsilon (b - a) \rightarrow 0 \text{ für } \varepsilon \rightarrow 0.
 \end{aligned}$$

□

32.8 Beispiel:

Kreisumfang: $c(t) = \begin{pmatrix} x(t) \\ y(t) \end{pmatrix} = \begin{pmatrix} r \cdot \cos t \\ r \cdot \sin t \end{pmatrix}, t \in [0, 2\pi]$

$$\begin{aligned}
 L(c) &= \int_0^{2\pi} |\dot{c}(t)| \, dt \\
 \dot{c}(t) &= \begin{pmatrix} \dot{x}(t) \\ \dot{y}(t) \end{pmatrix} = \begin{pmatrix} -r \cdot \sin t \\ r \cdot \cos t \end{pmatrix} \\
 |\dot{c}(t)| &= \sqrt{r^2 \sin^2 t + r^2 \cos^2 t} = \sqrt{r^2 (\sin^2 t + \cos^2 t)} = r \\
 \Rightarrow L(c) &= \int_0^{2\pi} r \, dt = 2\pi r. \quad \text{stimmt}
 \end{aligned}$$

Warum ist die Kurvenlänge wichtig?

32.9 Lemma: (Parametrisierungsinvarianz)

Die Länge einer C^1 -Kurve ist parametrisierungsinvariant.

Beweis:

Mit einem C^1 -Parameterwechsel $h: [\alpha, \beta] \rightarrow [a, b]$ gilt aufgrund der Substitui-

onsregel

$$\begin{aligned}
 L(c \circ h) &= \int_{\alpha}^{\beta} \left| \dot{c}(h(\tau)) \underbrace{h'(\tau)}_{>0} \right| d\tau && \text{(Kettenregel)} \\
 &= \int_{\tau=\alpha}^{\tau=\beta} |\dot{c}(h(\tau))| \cdot \underbrace{h'(\tau)}_{\frac{d}{d}t} dt \\
 &= \int_{t=a}^{t=b} |\dot{c}(t)| dt && \text{(Substitutionsregel)} \\
 &= L(c).
 \end{aligned}$$

□

32.10 Definition: (Bogenlängenfunktion)

Es sei $c : [a, b] \rightarrow \mathbb{R}^n$ eine C^1 -Kurve. Die Funktion

$$s(t) := \int_a^t |\dot{c}(\tau)| d\tau \quad t \in [a, b]$$

heißt Bogenlängenfunktion von c .

Bedeutung: Reparametrisiert man eine Kurve mit ihrer Bogenlänge (Bogenlängenparametrisierung), so vereinfachen sich viele Formeln.

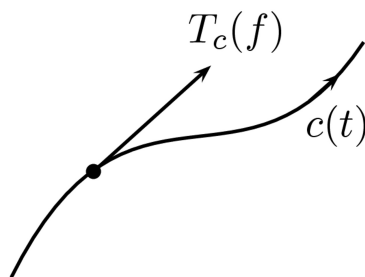
32.11 Definition: (Tangentenvektor, Tangenteneinheitsvektor)

Es sei $c : [a, b] \rightarrow \mathbb{R}^n$ eine C^1 -Kurve. Dann nennt man

$$\dot{c}(t) = (\dot{x}_1(t), \dots, \dot{x}_n(t))^T$$

auch Tangentenvektor (Geschwindigkeitsvektor) der Kurve c an der Stelle t .

Für $\dot{c}(t) \neq 0$ heißt $T_c(t) := \frac{\dot{c}(t)}{|\dot{c}(t)|}$ der Tangenteneinheitsvektor.



Bemerkung: Bei Bogenlängenparametrisierung (!) ist $|\dot{c}(t)| = 1$, d.h. $T_c(s) := \dot{c}(s)$ ist bereits Tangenteneinheitsvektor. Aus

$$1 = |\dot{c}(s)|^2 = (\dot{x}_1(s))^2 + (\dot{x}_2(s))^2 + \dots + (\dot{x}_n(s))^2$$

folgt durch Differentiation nach der Bogenlänge s :

$$0 = 2\dot{x}_1\ddot{x}_1 + 2\dot{x}_2\ddot{x}_2 + \dots + 2\dot{x}_n\ddot{x}_n = 2 \left\langle \begin{pmatrix} \dot{x}_1 \\ \vdots \\ \dot{x}_n \end{pmatrix}, \begin{pmatrix} \ddot{x}_1 \\ \vdots \\ \ddot{x}_n \end{pmatrix} \right\rangle$$

(Hierbei beschreibt $\langle \cdot, \cdot \rangle$ das Skalarprodukt in $\mathbb{R}^n$.)

Somit steht bei Bogenlängenparametrisierung der Beschleunigungsvektor $\ddot{c}(s) = \begin{pmatrix} \ddot{x}(s) \\ \ddot{y}(s) \end{pmatrix}$ senkrecht auf dem Geschwindigkeitsvektor $\dot{c}(s)$.

Man bezeichnet

$$N(s) := \frac{\ddot{c}(t)}{|\ddot{c}(t)|}$$

als Hauptnormalenvektor von c , und $\kappa(s) := |\ddot{c}(t)|$ als Krümmung von c .

32.12 Beispiel:

Krümmung eines Kreises mit Radius r

Kreis in Bogenlängenparametrisierung:

$$\begin{aligned} x(s) &= m_1 + r \cos \frac{s}{r} \\ y(s) &= m_2 + r \sin \frac{s}{r} \end{aligned}$$

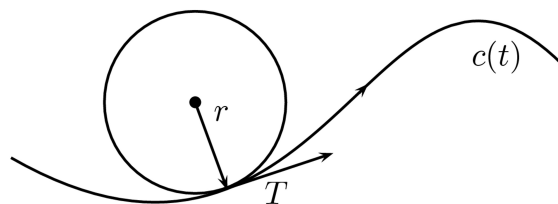
$$\begin{aligned} \dot{x}(s) &= -\sin \frac{s}{r} \\ \dot{y}(s) &= \cos \frac{s}{r} \end{aligned}$$

$$\begin{aligned} \ddot{x}(s) &= -\frac{1}{r} \cos \frac{s}{r} \\ \ddot{y}(s) &= -\frac{1}{r} \sin \frac{s}{r} \end{aligned}$$

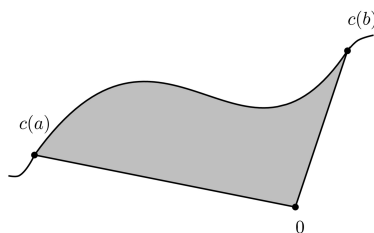
$$\kappa = \sqrt{\ddot{x}^2 + \ddot{y}^2} = \sqrt{\frac{1}{r^2} \cos^2 \frac{s}{r} + \frac{1}{r^2} \sin^2 \frac{s}{r}} = \frac{1}{r}$$

Bemerkung: Allgemein gilt: Die Krümmung gibt den reziproken Radius des Kreises an, der sich an der Stelle t an die Kurve $c(t)$ anschmiegt (Schmiegekreis, auch Krümmungskreis).

Der Schmiegekreis besitzt dieselbe Tangente und dieselbe Krümmung wie die Kurve.

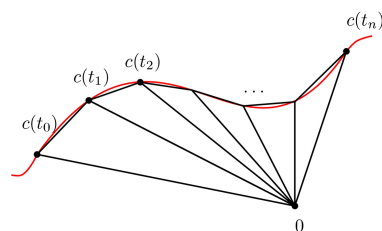


32.13 Die Sektorfläche



Ziel: Bestimmung des Flächeninhalts, der den Fahrstrahl von 0 nach $c(t)$ für $t \in [a, b]$ überstreicht.

Methode: approximiere durch Dreiecksflächen



Zerlegung

$$Z = \{t_0 = a < t_1 < \dots < t_n = b\}$$

Man kann zeigen (z.B. aus dem Schulunterricht bekannt):

Die Fläche des durch die Eckpunkte $0, c(t_i) = \begin{pmatrix} x_i \\ y_i \end{pmatrix}, c(t_{i+1}) = \begin{pmatrix} x_{i+1} \\ y_{i+1} \end{pmatrix}$ festgelegten Dreiecks beträgt

$$\frac{1}{2} |x_i y_{i+1} - x_{i+1} y_i|$$

Die Gesamtfläche beträgt daher

$$A(Z) = \frac{1}{2} \sum_{i=0}^{n-1} (x_i y_{i+1} - x_{i+1} y_i)$$

geeignete Orientierung der Dreiecke vorausgesetzt (so dass Beträge entfallen können). Deswegen wird im Folgenden ein orientierter Flächeninhalt eingeführt.

32.14 Definition: (Orientierter Flächeninhalt)

Der Fahrstrahl an die Kurve $c : [a, b] \rightarrow \mathbb{R}^2$ überstreicht den orientierten Flächeninhalt $F(c)$, wenn es zu jedem $\varepsilon > 0$ ein $\delta > 0$ gibt, so dass für jede Zerlegung Z von $[a, b]$ mit $\|Z\| \leq \delta$ gilt

$$|A(Z) - F(c)| \leq \varepsilon.$$

32.15 Satz: (Sektorformel von LEIBNIZ)

Sei $c : [a, b] \rightarrow \mathbb{R}^2$ eine stetig differenzierbare Kurve. Dann überstreicht der Ortsvektor $c(t)$ im Zeitintervall $t \in [a, b]$ die Fläche

$$F(c) = \frac{1}{2} \int_a^b (x\dot{y} - y\dot{x}) dt$$

Beweis:

Ähnlich zum Beweis von Satz 32.8 auf Seite 240

Nach dem Mittelwertsatz der Differentialrechnung gibt es in (t_i, t_{i+1}) Stellen τ_i und $\tilde{\tau}_i$ mit

$$\left. \begin{aligned} x_{i+1} - x_i &= \dot{x}(\tau_i)(t_{i+1} - t_i) \\ y_{i+1} - y_i &= \dot{y}(\tau_i)(t_{i+1} - t_i). \end{aligned} \right\} \quad (*)$$

Somit gilt für die Summen der Dreiecksflächen

$$\begin{aligned} A(Z) &= \frac{1}{2} \sum_{i=0}^{n-1} (x_i y_{i+1} - x_{i+1} y_i) \\ &= \frac{1}{2} \sum_{i=0}^{n-1} \left(\underbrace{x_i y_{i+1} - x_i y_i}_{x_i(y_{i+1} - y_i)} + \underbrace{x_i y_i - x_{i+1} y_i}_{-y_i(x_{i+1} - x_i)} \right) \\ &\stackrel{(*)}{=} \sum_{i=0}^{n-1} (x_i \dot{y}(\tilde{\tau}_i) - y_i \dot{x}(\tau_i)) (t_{i+1} - t_i) \end{aligned}$$

Wir vergleichen $A(Z)$ mit der Riemann-Summe

$$R(Z) = \frac{1}{2} \sum_{i=0}^{n-1} (x_i \dot{y}_i - y_i \dot{x}_i) (t_{i+1} - t_i).$$

Sei $\epsilon > 0$ gegeben und $\delta > 0$ so gewählt, dass gilt:

- (i) Für jede Zerlegung Z mit $\|Z\| \leq \delta$ ist

$$\left| R(Z) - \frac{1}{2} \int_a^b xy - yx \, dt \right| \leq \epsilon$$

- (ii) Für alle Paare $t, s \in [a, b]$ mit $|t - s| \leq \delta$ ist

$$\left. \begin{array}{l} |\dot{x}(t) - \dot{x}(s)| \leq \epsilon \\ |\dot{y}(t) - \dot{y}(s)| \leq \epsilon \end{array} \right\} \text{ Stetigkeit von } \dot{x}, \dot{y}$$

Es sei Z Zerlegung von $[a, b]$ mit $\|Z\| \leq \delta$. Ferner sei M obere Schranke für $|\dot{x}(t)|$ und $|\dot{y}(t)|$ auf $[a, b]$. Dann gilt:

$$\begin{aligned} |A(Z) - R(Z)| &\leq \frac{M}{2} \sum_{i=0}^{n-1} (|\dot{y}(\tilde{\tau}_i) - \dot{y}_i| + |\dot{x}(\tau_i) - \dot{x}_i|) (t_{i+1} - t_i) \\ &\leq \epsilon M \sum_{i=0}^{n-1} (t_{i+1} - t_i) = \epsilon M(b - a). \end{aligned}$$

Zusammen mit (i) ergibt sich

$$\begin{aligned} \left| A(Z) - \frac{1}{2} \int_a^b (xy - yx) \, dt \right| &\leq |A(Z) - R(Z)| + \left| R(Z) - \frac{1}{2} \int_a^b (xy - yx) \, dt \right| \\ &\leq \epsilon(M(b - a) + 1) \longrightarrow 0 \text{ für } \epsilon \rightarrow 0 \end{aligned}$$

Nach Def. 32.14 auf der vorherigen Seite folgt hieraus die Behauptung.

□

32.16 Beispiele

- (a) Der Fahrstrahl an den orientierten Kreisbogen

$$\begin{aligned} x &= r \cos t \\ y &= r \sin t \quad t \in [0, \varphi] \end{aligned}$$

überstreicht die orientierte Fläche

$$\frac{1}{2} \int_0^\varphi (x\dot{y} - y\dot{x}) \, dt = \frac{1}{2} \int_0^\varphi r \cos t \cdot r \cos t + r \sin t \cdot r \sin t \, dt = \frac{r^2}{2} \varphi.$$

Für $\phi = 2\pi$ ergibt sich die Kreisfläche πr^2 .

- (b) Der Fahrstrahl an den Zykloidenbogen

$$\begin{aligned} x &= t - \sin t \\ y &= 1 - \cos t \quad t \in [0, 2\pi] \end{aligned}$$

überstreicht die (orientierte) Fläche

$$\begin{aligned}\frac{1}{2} \int_0^{2\pi} ((t - \sin t) \sin t - (1 - \cos t)^2) \, dt &= \frac{1}{2} \int_0^{2\pi} t \sin t \sin^2 t - 1 + 2 \cos t - \cos^2 t \, dt \\ &= \frac{1}{2} [-t \cos t]_0^{2\pi} + \frac{1}{2} \int_0^{2\pi} 3 \cos t - 2 \, dt \\ &= \frac{1}{2} [-t \cos t + 3 \sin t - 2t]_0^{2\pi} = -3\pi.\end{aligned}$$

(Negativ, wegen mathematisch negativen Drehsinns)

Kapitel 33

Mehrfachintegrale

33.1 Motivation

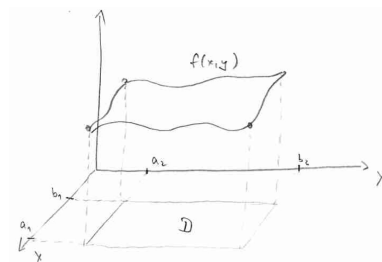
Funktionen können von mehreren Variablen abhängen.

Beispiel: $f(x, y) = x^2 - 3y$

Lassen sich für solche Funktionen Begriffe wie das Integral sinnvoll definieren?

Wie beschränken uns zunächst auf Funktion $f : D \rightarrow \mathbb{R}$ mit einem rechteckigen Definitionsbereich $D := [a_1, b_1] \times [a_2, b_2] \subset \mathbb{R}^2$.

Ziel ist die Berechnung des Volumens unterhalb des Graphen von f .



33.2 Definition: (Zerlegung, Feinheit, Flächeninhalt)

- (a) $Z = \{(x_0, x_1, \dots, x_n), (y_0, y_1, \dots, y_m)\}$ heißt Zerlegung des Rechtecks $D = [a_1, b_1] \times [a_2, b_2]$, falls

$$a_1 = x_0 < x_1 < \dots < x_n = b_1$$

$$a_2 = y_0 < y_1 < \dots < y_m = b_2$$

$\mathfrak{Z}$ ist die Menge der Zerlegungen von D , und

$$\|Z\| := \max_{i,j} \{|x_{i+1} - x_i|, |y_{j+1} - y_j|\}$$

die Feinheit einer Zerlegung $Z \in \mathfrak{Z}(D)$.

(b) Zu einer Zerlegung Z heißen die Mengen

$$\begin{aligned} Q_{ij} &:= [x_i, x_{i+1}] \times [y_j, y_{j+1}] \text{ die Teilquader von } Z, \text{ und} \\ \text{vol}(Q_{ij}) &:= (x_{i+1} - x_i) \cdot (y_{j+1} - y_j) \text{ ist der } \underline{\text{Flächeninhalt}} \text{ von } Q_{ij}. \end{aligned}$$

(c)

$$\begin{aligned} U_f(Z) &:= \sum_{i,j} \inf_{x \in Q_{ij}} (f(x)) \cdot \text{vol}(Q_{ij}) \\ O_f(Z) &:= \sum_{i,j} \sup_{x \in Q_{ij}} (f(x)) \cdot \text{vol}(Q_{ij}) \end{aligned}$$

heißen Riemann'sche Unter- bzw. Obersumme der Funktion f zur Zerlegung von Z .

Für beliebige Punkte $x_{ij} \in Q_{ij}$ heißt

$$R_f(Z) := \sum_{i,j} (f(x_{ij})) \cdot \text{vol}(Q_{ij})$$

eine Riemann'sche Summe zur Zerlegung Z .

33.3 Bemerkungen:

Analog zu 28.4 auf Seite 206 gelten folgende Aussagen:

(a) Die Riemann-Summe liegt zwischen Unter- und Obersumme:

$$U_f(Z) \leq R_f(Z) \leq O_f(Z).$$

(b) Entsteht eine Zerlegung Z_2 aus Z_1 durch Hinzunahme weiterer Zerlegungspunkte x_i und/oder y_j , so gilt:

$$\begin{aligned} U_f(Z_2) &\geq U_f(Z_1) \\ O_f(Z_2) &\leq O_f(Z_1). \end{aligned}$$

(c) Für beliebige Zerlegungen $Z_1, Z_2 \in \mathfrak{Z}$ gilt stets

$$U_f(Z_1) \leq O_f(Z_2)$$

Kann man mit diesen Unter-, Ober- und Riemann-Summen ein Integral definieren?

33.4 Definition: (*Unterintegral, Oberintegral, Riemann-Integral*)

(a) *Existieren die Grenzwerte*

$$\begin{aligned}\int \int_D f(x, y) \, dx \, dy &:= \sup\{U_f(Z) \mid Z \in \mathfrak{Z}(D)\} \\ \int \int_D^- f(x, y) \, dx \, dy &:= \inf\{O_f(Z) \mid Z \in \mathfrak{Z}(D)\}\end{aligned}$$

so heißen sie Riemann'sches Unter- bzw. Oberintegral.

(b) *$f(x, y)$ heißt (Riemann-integrierbar) über D , falls Unter- und Oberintegral übereinstimmen. Dann heißt*

$$\int \int_D f(x, y) \, dx \, dy := \int \int_D f(x, y) \, dx \, dy = \int \int_D^- f(x, y) \, dx \, dy$$

das (Riemann-)Integral von $f(x, y)$ über D .

Ähnlich wie im 1D-Fall zeigt man

33.5 Satz: (*Eigenschaften des Mehrfachintegrals*)

(a) Linearität

$$\int \int_D (\alpha f(x, y) + \beta g(x, y)) \, dx \, dy = \alpha \int \int_D f(x, y) \, dx \, dy + \beta \int \int_D g(x, y) \, dx \, dy$$

(b) Monotonie

Gilt $f(x, y) \leq g(x, y)$ für alle $(x, y) \in D$, so folgt

$$\int \int_D f(x, y) \, dx \, dy \leq \int \int_D g(x, y) \, dx \, dy.$$

(c) Nichtnegativität

$$f(x, y) \geq 0 \, \forall (x, y) \in D \rightarrow \int \int_D f(x, y) \, dx \, dy \geq 0$$

(d) Zusammensetzung von Integrationsbereichen

Sind D_1, D_2, D Rechtecke mit $D = D_1 \cup D_2$ und $\text{vol}(D_1 \cap D_2) = 0$, so ist $f(x, y)$ genau dann über D integrierbar, falls $f(x, y)$ über D_1 und über D_2 integrierbar ist. Dann gilt:

$$\int \int_D f(x, y) \, dx \, dy = \int \int_{D_1} f(x, y) \, dx \, dy + \int \int_{D_2} f(x, y) \, dx \, dy.$$

(e) **Riemann-Kriterium**

$f(x, y)$ ist genau dann über D integrierbar, falls gilt:

$$\forall \epsilon > 0 \exists Z \in \mathfrak{Z}(D) : O_f(Z) - U_f(Z) < \epsilon.$$

Schön wäre es, wenn sich die Berechnung von Doppelintegralen auf die Integration von Funktionen mit einer Variablen zurückführen ließe. Dass dies tatsächlich möglich ist, besagt der folgende Satz:

33.6 Satz: (Satz von Fubini)

Sei $D = [a_1, b_1] \times [a_2, b_2] \subset \mathbb{R}^2$ ein Rechteck. Ist $f : D \rightarrow \mathbb{R}$ integrierbar und existieren die Integrale

$$F(x) = \int_{a_2}^{b_2} f(x, y) \, dy \quad \text{bzw.} \quad G(y) = \int_{a_1}^{b_1} f(x, y) \, dx$$

für alle $x \in [a_1, b_1]$ bzw. $y \in [a_2, b_2]$, so gilt:

$$\begin{aligned} \int \int_D f(x, y) \, dx \, dy &= \int_{a_1}^{b_1} \left(\int_{a_2}^{b_2} f(x, y) \, dy \right) dx \\ &\text{bzw.} \\ \int \int_D f(x, y) \, dx \, dy &= \int_{a_2}^{b_2} \left(\int_{a_1}^{b_1} f(x, y) \, dx \right) dy \end{aligned}$$

Beweis:

Sei $Z = \{(x_0, x_1, \dots, x_n), (y_0, y_1, \dots, y_m)\}$ eine Zerlegung von D . Für beliebige $y \in [y_i, y_{i+1}]$, $\xi_i \in [x_i, x_{i+1}]$ gilt dann:

$$\inf_{(x,y) \in Q_{ij}} f(x, y) \leq f(\xi_i, y) \leq \sup_{(x,y) \in Q_{ij}} f(x, y).$$

Integration bzgl. y liefert

$$\inf_{(x,y) \in Q_{ij}} f(x, y) \leq \int_{y_j}^{y_{j+1}} f(\xi_i, y) \, dy \leq \sup_{Q_{ij}} (f(x, y))(y_{i+1} - y_i)$$

Multiplikation mit $(x_{i+1} - x_i)$ und Summation über i, j ergibt

$$U_f(Z) \leq \sum_{i=0}^{n-1} \left(\int_{a_2}^{b_2} f(\xi_i, y) \, dy \right) (x_{i+1} - x_i) \leq O_f(Z)$$

Das ist eine Abschätzung für eine beliebige Riemann-Summe von

$$F(x) = \int_{a_2}^{b_2} f(x, y) \, dy \text{ zur Zerlegung } Z_x = \{x_0, \dots, x_n\} \text{ von } [a_1, b_1].$$

Insbesondere folgt hieraus:

$$U_f(Z) \leq U_f(Z_x) \leq O_f(Z_x) \leq O_f(Z).$$

Für $\|Z\| \rightarrow 0$ erhält man die Behauptung.

□

33.7 Beispiel:

Sei $D = [0, 1] \times [0, 2]$ und $f(x, y) = 2 - xy$

$$\begin{aligned}\iint_D f(x, y) \, dx \, dy &= \int_{y=0}^{y=2} \left(\int_{x=0}^{x=1} f(x, y) \, dx \right) dy \\&= \int_{y=0}^{y=2} \left[2x - \frac{1}{2}x^2 y \right]_{x=0}^{x=1} dy = \int_0^2 \left(2 - \frac{1}{2}y \right) dy \\&= \left[2y - \frac{1}{4}y^2 \right]_0^2 \\&= 4 - \frac{1}{4} \cdot 4 = 3.\end{aligned}$$

33.8 Bemerkungen:

- (a) Das Doppelintegral $\iint_D f(x, y) \, dx \, dy$ beschreibt das Volumen zwischen dem Graphen von $f(x, y)$ und der $x - y$ -Ebene.
- (b) Die vorangegangenen Betrachtungen lassen sich direkt auf Dreifachintegrale usw. verallgemeinern.

Teil IV

Lineare Algebra

Kapitel 34

Vektorräume

34.1 Motivation

- Im $\mathbb{R}^2$ und $\mathbb{R}^3$ kann man Vektoren addieren und mit einem Skalar multiplizieren.
- Wir wollen dieses Grundkonzept algebraisch formalisieren, um es auch auf andere Situationen anwenden zu können.
- Wichtig z.B. in der Robotik, der Codierungstheorie, in Computergrafik und Computer Vision.

34.2 Definition: (K -Vektorraum)

Sei K ein Körper. Ein K -Vektorraum ist eine Menge V , auf der eine Verknüpfung $+$: $V \times V \rightarrow V$ und eine skalare Multiplikation $\cdot$: $K \times V \rightarrow V$ definiert sind mit

- (a) $(V, +)$ ist eine kommutative Gruppe.
- (b) $\lambda \cdot (\mu \cdot v) = (\lambda\mu) \cdot v \quad \forall \lambda, \mu \in K, \forall v \in V$
- (c) $1 \cdot v = v \quad \forall v \in V$
- (d) $\lambda \cdot (v + w) = \lambda \cdot v + \lambda \cdot w \quad \forall \lambda \in K, \forall v, w \in V$
- (e) $(\lambda + \mu) \cdot v = \lambda \cdot v + \mu \cdot v \quad \forall \lambda, \mu \in K, \forall v \in V$

Die Elemente von V heißen Vektoren, Elemente aus K heißen Skalare. Ist $K = \mathbb{R}$ oder $K = \mathbb{C}$, sprechen wir von einem reellen bzw. komplexen Vektorraum.

Bemerkung: Für Skalare verwenden wir meist kleine griechische Buchstaben wie λ, μ, ν . Vektoren bezeichnen wir mit kleinen lateinischen Buchstaben wie

u, v, w . Nur wenn Verwechslungsgefahr besteht, verwenden wir Pfeile auf Vektoren, z.B. um den Nullvektor $\vec{0} = \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix}$ von der skalaren Null 0 zu unterscheiden.

34.3 Satz: (Rechenregeln für Vektorräume)

In einem K -Vektorraum V gilt:

- (a) $\lambda \cdot \vec{0} = \vec{0} \quad \forall \lambda \in K.$
- (b) $0 \cdot v = \vec{0} \quad \forall v \in V.$
- (c) $(-1)v = -v \quad \forall v \in V.$

Beweis:

(von a)

$$\lambda \cdot \vec{0} \stackrel{\vec{0}=\vec{0}+\vec{0}}{=} \lambda \cdot (\vec{0} + \vec{0}) \stackrel{34.2(d)}{=} \lambda \cdot \vec{0} + \lambda \cdot \vec{0}$$

Subtrahiert man auf beiden Seiten $\lambda \cdot \vec{0}$, folgt $\vec{0} = \lambda \cdot \vec{0}$.

□

34.4 Beispiele

- (a) $\mathbb{R}^n$ ist ein $\mathbb{R}$ -Vektorraum für alle $n \in \mathbb{N}$.
- (b) $\mathbb{C}^n$ ist ein $\mathbb{C}$ -Vektorraum für alle $n \in \mathbb{N}$.
- (c) $\mathbb{R}$ ist ein $\mathbb{Q}$ -Vektorraum.
- (d) Für einen Körper K ist K^n für alle $n \in \mathbb{N}$ ein K -Vektorraum.
- (e) Insbesondere ist $\mathbb{Z}_2^n$ ein Vektorraum über dem Körper $\mathbb{Z}_2$. Er besteht aus allen n -Tupeln von Nullen und Einsen. Beispielsweise lassen sich Integer-Werte in Binärdarstellung als Element des Vektorraums $\mathbb{Z}_2^{32}$ darstellen. Die Vektorräume $\mathbb{Z}_2^n$ spielen eine große Rolle in der Codierungstheorie. Lineare Codes verwenden Teilmengen des $\mathbb{Z}_2^n$, bei denen beim Auftreten eines Übertragungsfehlers mit hoher Wahrscheinlichkeit ein Element außerhalb der Teilmenge entsteht (vgl. 11.3 auf Seite 96).
- (f) Funktionsräume sind wichtige Vektorräume. Definiert man auf

$$\mathcal{F} := \{f \mid \mathbb{R} \rightarrow \mathbb{R}\}$$

eine Vektoraddition und eine skalare Multiplikation durch

$$\begin{aligned}(f + g)(x) &:= f(x) + g(x) & f, g \in \mathcal{F} \\ (\lambda \cdot g)(x) &:= \lambda \cdot g(x) & \lambda \in \mathbb{R}, \forall g \in \mathcal{F}\end{aligned}$$

so ist $\mathcal{F}$ ein $\mathbb{R}$ -Vektorraum. Der Nullvektor $\vec{0}$ ist die Funktion $f(x) = 0 \quad \forall x \in \mathbb{R}$.

- (g) Die Polynome $K[x]$ über einem Körper K bilden einen K -Vektorraum, wenn man als skalare Multiplikation definiert:

$$\lambda \sum_{k=0}^n a_k x^k := \sum_{k=0}^n \lambda \cdot a_k x^k \quad \forall \lambda \in K.$$

Zu Gruppen und Ringe haben wir in 8.5 auf Seite 72 und 9.5 auf Seite 80 Untergruppen und Unterringe definiert. Ähnliches ist auch für Vektorräume möglich.

34.5 Definition: (Unterraum, Untervektorraum, Teilraum)

Sei V ein K -Vektorraum und $U \subset V$. Ist U mit den Verknüpfungen von V selbst wieder ein K -Vektorraum, so heißt U Unterraum (Untervektorraum, Teilraum) von V .

34.6 Satz: (Unterraumkriterium)

Sei V ein K -Vektorraum und $U \subset V$ eine nichtleere Teilmenge. Dann ist U genau dann ein Unterraum von V , wenn gilt:

- (a) $u + v \in U \quad \forall u, v \in U$
(b) $\lambda \cdot u \in U \quad \forall u \in U, \forall \lambda \in K$.

Beweis:

„ $\Rightarrow$ “ Offensichtlich, da U selbst ein Vektorraum ist.

„ $\Leftarrow$ “ Nach 8.5 auf Seite 72 ist U Untergruppe von V , denn

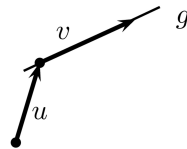
$$\begin{aligned}u + v &\in U & \forall u, v &\in U \\ -U &\stackrel{\text{34.3 auf dervorherigenSeite(c)}}{=} (-1) \cdot u &\stackrel{(b)}{\in} U\end{aligned}$$

Da V kommutativ ist, ist U eine kommutative Gruppe, d.h. 34.2 (a) ist erfüllt.

Aus (b) folgt, dass $\cdot$ eine Abbildung von $K \times U$ nach U ist. Die Eigenschaften (b)-(e) der Vektorraumdefinition 34.2 übertragen sich von der Vektorraumeigenschaft von V .

34.7 Beispiele

- (a) Ein K -Vektorraum hat die trivialen Unterräume $\{0\}$ und V .
- (b) Lineare Codes sind Unterräume im $\mathbb{Z}_2^n$.



- (c) Sind Geraden Unterräume des $\mathbb{R}^2$? Ein Unterraum muss stets den Nullvektor enthalten. Daher bilden nur die Ursprungsgeraden $\{\lambda v \mid \lambda \in \mathbb{R}\}$ Unterräume.
- (d) Verallgemeinerung von (c):
In einem K -Vektorraum V bilden die Mengen $\{\lambda v \mid \lambda \in K\}$ Unterräume.

Das letzte Beispiel läßt sich noch weiter verallgemeinern:

34.8 Definition: (Linearkombination, Erzeugnis)

Sei V ein K -Vektorraum. Ferner seien $u_1, \dots, u_n \in V$ und $\lambda_1, \dots, \lambda_n \in K$. Dann nennt man $\sum_{k=1}^n \lambda_k u_k$ eine Linearkombination von $u_1, \dots, u_n$. Die Menge aller Linearkombinationen bildet das Erzeugnis $\text{Span}(u_1, \dots, u_n)$, und man nennt $\{u_1, \dots, u_n\}$ das Erzeugendensystem von $\text{Span}(u_1, \dots, u_n)$.

34.9 Satz: (Erzeugnis als Unterraum)

Sei V ein K -Vektorraum und seien $u_1, \dots, u_n \in V$. Dann bildet $\text{Span}(u_1, \dots, u_n)$ einen Unterraum von V .

Beweis:

Wir wenden das Unterraumkriterium an:

- (a) Seien $v, w \in \text{Span}(u_1, \dots, u_n) \Rightarrow \exists \lambda_1, \dots, \lambda_n$ und $\mu_1, \dots, \mu_n \in K$ mit

$$\begin{aligned}
 v &= \sum_{i=1}^n \lambda_i u_i, & w &= \sum_{i=1}^n \mu_i u_i \\
 \Rightarrow v + w &= \sum_{i=1}^n \lambda_i u_i + \sum_{i=1}^n \mu_i u_i \\
 &\stackrel{34.2(e)}{=} \sum_{i=1}^n (\lambda_i + \mu_i) u_i \in \text{Span}(u_1, \dots, u_n).
 \end{aligned}$$

(b) Sei $\mu \in K$ und $u \in \text{Span}(u_1, \dots, u_n)$.

$$\begin{aligned} \Rightarrow \exists \lambda_1, \dots, \lambda_n \in K : u &= \sum_{i=1}^n \lambda_i u_i \\ \Rightarrow \mu u &= \mu \sum_{i=1}^n \underbrace{\lambda_i u_i}_{\in V} \\ &\stackrel{34.2(d)}{=} \sum_{i=1}^n \mu(\lambda_i u_i) \\ &\stackrel{34.2(b)}{=} \sum_{i=1}^n \underbrace{(\mu \lambda_i)}_{\in K} u_i \in \text{Span}(u_1, \dots, u_n). \end{aligned}$$

□

34.10 Beispiel:

Im $\mathbb{R}^3$ bildet die Menge aller Ebenen durch den Ursprung einen Unterraum.

34.11 Lineare Abhängigkeit

Offenbar ist $\text{Span} \left(\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \right) = \mathbb{R}^3 = \text{Span} \left(\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} \right)$,
d.h. $\begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}$ läßt sich als Linearkombination von $\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$ darstellen.

Definition:

linear abhängig, linear unabhängig Ein Vektor u heißt linear abhängig von $u_1, \dots, u_n$, wenn es $\lambda_1, \dots, \lambda_n \in K$ gibt mit $\sum_{i=1}^n \lambda_i u_i$.

Eine Menge von Vektoren $u_1, \dots, u_n$, bei denen sich keiner der Vektoren als Linearkombination der anderen ausdrücken lässt, heißt linear unabhängig.

Gibt es ein einfaches Kriterium, mit dem man lineare Unabhängigkeit nachweisen kann?

34.12 Satz: (Kriterium für lineare Unabhängigkeit)

Sei V ein K -Vektorraum. Die Vektoren $v_1, \dots, v_n \in V$ sind genau dann linear unabhängig, wenn für jede Linearkombination mit $\sum_{i=1}^n \lambda_i v_i = \vec{0}$ gilt:

$$\lambda_1 = \dots = \lambda_n = 0.$$

Beweis:

„ $\Rightarrow$ “ : Anm.: Es gibt eine Linearkombination, in der ein $\lambda_k \neq 0$ existiert und $\sum_{i=1}^n \lambda_i v_i = \vec{0}$ gilt.

$$\begin{aligned} \Rightarrow -\lambda_k v_k &= \sum_{\substack{i=1 \\ i \neq k}}^n \lambda_i v_i \\ \Rightarrow v_k &= \sum_{\substack{i=1 \\ i \neq k}}^n \underbrace{\frac{\lambda_i}{-\lambda_k}}_{\in K} v_i \end{aligned}$$

$\Rightarrow v_k$ ist linear abhängig von $\{v_1, \dots, v_n\} \setminus \{v_k\}$.

„ $\Leftarrow$ “ : Seien $v_1, \dots, v_n$ linear abhängig, d.h. es gibt ein v_k mit $v_k = \sum_{\substack{i=1 \\ i \neq k}}^n \lambda_i v_i$.

$$\text{Setze } \lambda_k := -1 \Rightarrow 0 = \sum_{i=1}^n \lambda_i v_i.$$

□

Macht lineare Unabhängigkeit auch bei unendlichen vielen Vektoren Sinn?

34.13 Definition: (Lineare Unabhängigkeit)

Ein unendliches System B von Vektoren heißt linear unabhängig, wenn jede endliche Auswahl von Vektoren aus B linear unabhängig ist.

34.14 Beispiel:

Betrachte $\mathbb{R}[x]$, d.h. die Menge aller Polynome mit Koeffizienten aus $\mathbb{R}$. Nach 34.4 (g) bildet $\mathbb{R}[x]$ einen $\mathbb{R}$ -Vektorraum.

Wir zeigen, dass das unendliche System $B = \{1, x, x^2, \dots\}$ linear unabhängig in $\mathbb{R}[x]$ ist.

Anm.: Dies trifft nicht zu. Dann existiert eine endliche Teilmenge $\{x^{m_1}, x^{m_2}, \dots, x^{m_n}\} \subset B$, die linear abhängig ist.

$\Rightarrow$ Es gibt eine Linearkombination

$$\sum_{i=1}^n \lambda_i v^{m_i} = 0 \quad (*)$$

mit einem $\lambda_k \neq 0$. Auf der linken Seite von $(*)$ steht ein Polynom, das nur endlich viele Nullstellen hat, rechts steht das Nullpolynom mit unendlich vielen Nullstellen. `

□

Der Begriff der linearen Unabhängigkeit erlaubt uns ein minimales Erzeugendensystem zu finden.

34.15 Definition: (*Basis*)

Sei V ein Vektorraum. Eine Teilmenge $B \subset V$ heißt Basis von V , falls gilt:

- (a) $\text{Span}(B) = V$
- (b) B ist linear unabhängig.

34.16 Beispiele

- (a) $\left\{ \begin{pmatrix} 2 \\ 3 \end{pmatrix}, \begin{pmatrix} 3 \\ 4 \end{pmatrix} \right\}$ ist eine Basis des $\mathbb{R}^2$, denn:

- (i) Sei $\begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{R}^2$. Wir suchen λ, μ mit

$$\lambda \begin{pmatrix} 2 \\ 3 \end{pmatrix} + \mu \begin{pmatrix} 3 \\ 4 \end{pmatrix} = \begin{pmatrix} x \\ y \end{pmatrix}$$

Das Gleichungssystem

$$\begin{aligned} 2\lambda + 3\mu &= x \\ 3\lambda + 4\mu &= y \end{aligned}$$

hat die Lösung $\lambda = -4x + 3y, \mu = 3x - 2y$. (*)

$$\Rightarrow \mathbb{R}^2 = \text{Span} \left(\begin{pmatrix} 2 \\ 3 \end{pmatrix}, \begin{pmatrix} 3 \\ 4 \end{pmatrix} \right)$$

- (ii) Sei

$$\lambda \begin{pmatrix} 2 \\ 3 \end{pmatrix} + \mu \begin{pmatrix} 3 \\ 4 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}.$$

Mit (*) folgt:

$$\lambda = -4 \cdot 0 + 3 \cdot 0 = 0$$

$$\mu = 3 \cdot 0 - 2 \cdot 0 = 0$$

$\Rightarrow \begin{pmatrix} 2 \\ 3 \end{pmatrix}, \begin{pmatrix} 3 \\ 4 \end{pmatrix}$ sind linear unabhängig.

$$(b) \left\{ \begin{pmatrix} 1 \\ 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \dots, \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \\ 1 \end{pmatrix} \right\}$$

bildet eine Basis des $\mathbb{R}^n$, die so genannte Standardbasis.

(c) Es gibt auch Vektorräume mit unendlichen Basen. Beispielsweise ist $\{1, x, x^2, \dots\}$ eine unendliche Basis von $\mathbb{R}[x]$:

(i) $\{1, x, x^2, \dots\}$ ist linear unabhängig nach 34.14 auf Seite 258.

(ii) Jedes Polynom aus $\mathbb{R}[x]$ lässt sich als Linearkombination von Elementen aus $\{1, x, x^2, \dots\}$ darstellen.

Hat jeder Vektorraum eine Basis? Man kann zeigen:

34.17 Satz: (Existenz einer Basis)

Jeder Vektorraum $V \neq \{\vec{0}\}$ hat eine Basis.

Offenbar sind Basen nicht eindeutig: So sind z.B.

$$\left\{ \begin{pmatrix} 2 \\ 3 \end{pmatrix}, \begin{pmatrix} 3 \\ 4 \end{pmatrix} \right\} \text{ und } \left\{ \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right\} \text{ Basen des } \mathbb{R}^2.$$

Insbesondere kann man Basisvektoren austauschen:

34.18 Satz: (Austauschsatz)

Sei V ein Vektorraum mit einer endlichen Basis $B := \{b_1, \dots, b_n\}$. Ferner sei $v \in V$ und $v \neq \vec{0}$. Dann existiert ein b_k , so dass $\{b_1, \dots, b_{k-1}, v, b_{k+1}, \dots, b_n\}$ eine Basis von V ist.

Beweis:

Da B Basis $\Rightarrow \exists \lambda_1, \dots, \lambda_n$ mit

$$v = \sum_{i=1}^n \lambda_i b_i. \quad (*)$$

Da $v \neq 0$ ist mindestens ein $\lambda_i \neq 0$. Sei o.B.d.A. $\lambda_1 \neq 0$.

Wir zeigen, dass dann $\{v, b_2, b_3, \dots, b_n\}$ eine Basis von V ist.

(i) $\text{Span}(v, b_2, b_3, \dots, b_n) = V$, denn:

Sei $u \in V$. Da B Basis ist, existiert $\mu_1, \dots, \mu_n$ mit

$$u = \sum_{k=1}^n \mu_k b_k. \quad (**)$$

Aus (*) und $\lambda_1 \neq 0$ folgt:

$$b_1 = \frac{1}{\lambda_1} v_1 - \sum_{i=2}^n \frac{\lambda_i}{\lambda_1} b_i.$$

Einsetzen in (**) zeigt, dass u als Linearkombination von $v, b_2, b_3, \dots, b_n$ geschrieben werden kann.

(a) $v, b_2, b_3, \dots, b_n$ ist linear unabhängig, denn:

Sei $\mu_1 \cdot v + \mu_2 \cdot b_2 + \mu_3 \cdot b_3 + \dots + \mu_n \cdot b_n = 0$.

Mit (*):

$$\begin{aligned} 0 &= \mu_1 \sum_{i=1}^n \lambda_i b_i + \sum_{i=2}^n \mu_i b_i \\ &= \mu_1 \lambda_1 b_1 + \sum_{i=2}^n (\mu_1 \lambda_i + \mu_i) b_i \end{aligned}$$

Da $\{b_1, \dots, b_n\}$ Basis, sind alle Koeffizienten 0:

$$\begin{aligned} \mu_1 \lambda_1 &= 0 \xrightarrow{\lambda_1 \neq 0} \mu_1 = 0 \\ \mu_1 \lambda_i + \mu_i &= 0 \xrightarrow{\mu_1 \neq 0} \mu_i = 0 \quad (i = 2, \dots, n). \end{aligned}$$

□

Hat ein Vektorraum eine eindeutige Zahl an Basisvektoren?

34.19 Satz: (Eindeutigkeit der Zahl der Basisvektoren)

Hat ein Vektorraum V eine endliche Basis von n Vektoren, so hat jede Basis von V ebenfalls n Vektoren.

Beweis:

Seien $B = \{b_1, \dots, b_n\}$ und $C = \{c_1, \dots, c_m\}$ Basen.

- (a) Anm.: $n > m$. Tausche m der Vektoren aus B durch C aus. Da C eine Basis ist, sind die nicht ausgetauschten Vektoren aus B als Linearkombination von $c_1, \dots, c_m$ darstellbar.
Dies widerspricht der Basiseigenschaft.

(b) Anm.: $n < m$. Widerspruch wird analog konstruiert.

□

Satz 34.19 auf der vorherigen Seite motiviert:

34.20 Definition: (*Dimension*)

Sei $B = \{b_1, \dots, b_n\}$ eine Basis des Vektorraums V . Dann heißt die Zahl $n := \dim V$ die Dimension von V . Der Nullraum $\{\vec{0}\}$ habe die Dimension 0

Mann kann zeigen:

34.21 Satz: (*Basiskriterium linear unabhängiger Mengen*)

In einem Vektorraum der Dimension n ist jede linear unabhängige Menge mit n Vektoren eine Basis.

Bemerkung: Insbesondere gibt es in einem n -dimensionalen Vektorraum keine Basen mit $n - 1$ oder $n + 1$ Elementen. Dies würde Satz 34.19 auf der vorherigen Seite widersprechen.

Kapitel 35

Lineare Abbildungen

35.1 Motivation

- Nachdem wir wichtige Eigenschaften von Vektorräumen kennen, macht es Sinn zu untersuchen, wie Abbildungen zwischen Vektorräumen aussehen können. Lineare Abbildungen sind die wichtigsten Abbildungen zwischen Vektorräumen.
- Der Basisbegriff liefert ein wichtiges Werkzeug zur Beschreibung linearer Abbildungen.

35.2 Definition: (lineare Abbildung, Vektorraumhomomorphismus)

Seien U, V K -Vektorräume. Eine Abbildung $f : U \rightarrow V$ heißt lineare Abbildung (Vektorraumhomomorphismus), falls gilt:

- (a) $f(u + v) = f(u) + f(v) \quad \forall u, v \in U$
- (b) $f(\lambda \cdot u) = \lambda f(u) \quad \forall \lambda \in K, \forall u \in U$

U und V heißen isomorph, wenn es eine bijektive lineare Abbildung $f : U \rightarrow V$ gibt. Wir schreiben hierfür $U \simeq V$.

Bemerkung: Ähnlich wie bei Gruppenhomomorphismen (vgl 8.1 auf Seite 69 (b)) überführt ein Vektorraumhomomorphismus die Verknüpfungen in U (skalare Multiplikation) in Verknüpfungen in V .

Die Bedingungen (a) und (b) fasst man oft in eine einzige Bedingung zusammen:

$$f(\lambda u + \mu v) = \lambda f(u) + \mu f(v) \quad \forall \lambda, \mu \in K, \forall u, v \in U.$$

D.h. Linearkombinationen in U werden in Linearkombinationen in V überführt.

35.3 Beispiele

(a) $f : \mathbb{R}^3 \rightarrow \mathbb{R}^2$:

$$\begin{aligned} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} &\mapsto \begin{pmatrix} 2x_1 + x_3 \\ -x_2 \end{pmatrix} \text{ ist linear:} \\ f\left(\lambda \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} + \mu \begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix}\right) &= f\left(\begin{pmatrix} \lambda x_1 + \mu y_1 \\ \lambda x_2 + \mu y_2 \\ \lambda x_3 + \mu y_3 \end{pmatrix}\right) \\ &= \begin{pmatrix} 2 \cdot (\lambda x_1 + \mu y_1) + \lambda x_3 + \mu y_3 \\ -\lambda x_2 - \mu y_2 \end{pmatrix} \\ &= \lambda \begin{pmatrix} 2x_1 + x_3 \\ -x_2 \end{pmatrix} + \mu \begin{pmatrix} 2y_1 + y_3 \\ -y_2 \end{pmatrix} \\ &= \lambda f\left(\begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix}\right) + \mu f\left(\begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix}\right) \\ &\quad \forall \lambda, \mu \in \mathbb{R}, \forall \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix}, \begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix} \in \mathbb{R}^3. \end{aligned}$$

(b) $f : \mathbb{R} \rightarrow \mathbb{R} \quad x \mapsto x + 1$ ist keine lineare Abbildung, denn

$$1 = f(0 + 0) \neq f(0) + f(0) = 1 + 1.$$

(c) $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2 : \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \mapsto \begin{pmatrix} x_1 \\ x_1 x_2 \end{pmatrix}$ ist ebenfalls nichtlinear:

$$\begin{aligned} f\left(2 \cdot \begin{pmatrix} 1 \\ 1 \end{pmatrix}\right) &= \begin{pmatrix} 2 \\ 4 \end{pmatrix}, \\ 2 \cdot f\left(\begin{pmatrix} 1 \\ 1 \end{pmatrix}\right) &= 2 \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 2 \\ 2 \end{pmatrix}. \end{aligned}$$

35.4 Bild, Kern, Monomorphismus,...

In Analogie zu 8.16 auf Seite 76 und 8.17 auf Seite 77 gibt es auch für Vektorräume Monomorphismen, Epimorphismen, Isomorphismen, Endomorphismen und Automorphismen. Ferner ist für $f : U \rightarrow V$:

$$\begin{aligned} \text{Im}(f) &:= \{f(u) \mid u \in U\} && \text{das \underline{Bild} von } f. \\ \ker(f) &:= \{u \in U \mid f(u) = 0\} && \text{den \underline{Kern} von } f. \end{aligned}$$

Für lineare Abbildungen lassen sich einige wichtige Eigenschaften zeigen:

35.5 Satz: (Eigenschaften linearer Abbildungen)

- a) Ist $f : U \rightarrow V$ ein Isomorphismus, so ist auch die Umkehrabbildung f^{-1} linear.
- b) Die lineare Abbildung $f : U \rightarrow V$ ist genau dann injektiv, wenn $\ker(f) = \{0\}$ ist.
- c) Ist $f : U \rightarrow V$ linear, so ist $\ker(f)$ ein Unterraum von U und $\operatorname{Im}(f)$ ein Unterraum von V .
- d) $\dim \ker(f) + \dim \operatorname{Im}(f) = \dim U$.

35.6 Beispiel:

$$f : \mathbb{R}^3 \rightarrow \mathbb{R}^2 : \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \mapsto \begin{pmatrix} x_1 - x_3 \\ 0 \end{pmatrix}$$

$\ker(f) = \left\{ \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \in \mathbb{R}^3 \mid x_1 = x_3 \right\}$ ist eine Ursprungsebene im $\mathbb{R}^3$ (und damit ein 2-dimensionaler Unterraum des Unterraum des $\mathbb{R}^3$).

$\operatorname{Im}(f) = \left\{ \begin{pmatrix} x \\ 0 \end{pmatrix} \in \mathbb{R}^2 \mid x \in \mathbb{R} \right\}$ ist eine Ursprungsgerade im $\mathbb{R}^2$ (und damit ein 1-dimensionaler Unterraum des $\mathbb{R}^2$).

$$\dim \ker(f) + \dim \operatorname{Im}(f) = 2 + 1 = 3 = \dim(\mathbb{R}^3).$$

Welche Rolle spielen Basen bei der Beschreibung linearer Abbildungen?
Hierzu betrachten wir zunächst Basisdarstellungen von Vektoren.

35.7 Satz: (Eindeutigkeit der Darstellung in einer festen Basis)

Sei $B = \{b_1, \dots, b_n\}$ eine Basis des K -Vektorraums V . Dann gibt es zu jedem Vektor $v \in V$ eindeutig(!) bestimmte Elemente $x_1, \dots, x_n \in K$ mit $v = \sum_{i=1}^n x_i b_i$.

Bemerkung: Diese x_i heißen Koordinaten von v bzgl. B . Wir schreiben

$$v = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}_B.$$

Beweis:

Die Existenz der Darstellung ist klar, da $V = \operatorname{Span}(B)$. Zu zeigen ist also nur

die Eindeutigkeit. Sei $\sum_{i=1}^n x_i b_i$ und $v = \sum_{i=1}^n y_i b_i$.

$$\Rightarrow \vec{0} = v - v = \sum_{i=1}^n (x_i - y_i) b_i$$

$$\Rightarrow x_i = y_i \quad (i = 1, \dots, n), \text{ da } b_i \text{ linear unabhängig.}$$

□

Selbstverständlich liefern unterschiedliche Basen auch unterschiedliche Koordinatendarstellungen eines Vektors. Wie kann man diese Darstellung in einander umrechnen?

35.8 Umrechnung von Koordinatendarstellungen

Beispiel: Im $\mathbb{R}^2$ sei eine Basis $B = \{b_1, b_2\}$ gegeben. Bzgl. B habe ein Vektor v die Darstellung $v = \begin{pmatrix} x \\ y \end{pmatrix}_B$.

Wir betrachten die neue Basis $C = \{c_1, c_2\}$ mit $c_1 = \begin{pmatrix} 2 \\ 2 \end{pmatrix}_B$, $c_2 = \begin{pmatrix} 1 \\ 2 \end{pmatrix}_B$.
Wie lautet die Darstellung von v in C ?

$$\begin{aligned} \text{Ansatz: } v &= \begin{pmatrix} \lambda \\ \mu \end{pmatrix}_C \\ \begin{pmatrix} x \\ y \end{pmatrix}_B &= v = \begin{pmatrix} \lambda \\ \mu \end{pmatrix}_C = \lambda c_1 + \mu c_2 \\ &= \lambda \begin{pmatrix} 2 \\ 2 \end{pmatrix}_B + \mu \begin{pmatrix} 1 \\ 2 \end{pmatrix}_B \\ &= \begin{pmatrix} 2\lambda + \mu \\ 2\lambda + 2\mu \end{pmatrix}_B. \end{aligned}$$

Die Bestimmungsgleichungen

$$\begin{aligned} 2\lambda + \mu &= x \\ 2\lambda + 2\mu &= y \end{aligned}$$

haben die Lösung

$$\begin{aligned} \lambda &= x - \frac{1}{2}y \\ \mu &= -x + y \end{aligned}$$

Beispielsweise ist

$$\begin{pmatrix} 5 \\ 2 \end{pmatrix}_B = \begin{pmatrix} 5 - \frac{1}{2} \cdot 2 \\ -5 + 2 \end{pmatrix}_C = \begin{pmatrix} 4 \\ -3 \end{pmatrix}_C.$$

Basen ermöglichen es, eine lineare Abbildung durch wenige Daten zu beschreiben: Es genügt zu wissen, was mit den Basisvektoren passiert.

35.9 Satz: (*Charakterisierung einer linearen Abbildung durch ihre Wirkung auf die Basis*)

Seien U, V K -Vektorräume und $b_1, \dots, b_n$ sei eine Basis von U . Ferner seien $v_1, \dots, v_n \in V$. Dann gibt es genau eine lineare Abbildung $f : U \rightarrow V$ mit $f(b_i) = v_i$ ($i = 1, \dots, n$).

Beweis:

Sei $u \in U$ und $u = \sum_{i=1}^n x_i b_i$. Setze $f(u) := \sum_{i=1}^n x_i v_i$.

Man prüft leicht nach, dass f linear ist und $f(b_i) = v_i$ für $i = 1, \dots, n$.

Zum Nachweis der Eindeutigkeit nehmen wir an, dass g eine weitere lineare Abbildung mit $g(b_i) = v_i \ \forall i$ ist.

$$\begin{aligned} \Rightarrow g(u) &= g\left(\sum_{i=1}^n x_i b_i\right) \\ &\stackrel{\text{lin.}}{=} \sum_{i=1}^n x_i g(b_i) = \sum_{i=1}^n x_i v_i \\ &= \sum_{i=1}^n x_i f(b_i) \\ &\stackrel{\text{lin.}}{=} f\left(\sum_{i=1}^n x_i b_i\right) = f(u) \end{aligned}$$

Somit stimmen f und g überall überein.

□

Als weiteres wichtiges Resultat, das auf Basen beruht, kann man zeigen:

35.10 Satz: (*Isomorphie endlichdimensionaler Vektorräume*)

Seien U, V endlichdimensionale K -Vektorräume. Dann sind U und V genau dann isomorph, wenn sie die selbe Dimension haben.

$$U \simeq V \Leftrightarrow \dim U = \dim V.$$

Bemerkung: Dieser Satz besagt z.B., dass es im Wesentlichen nur einen einzigen n -dimensionalen Vektorraum über $\mathbb{R}$ gibt: den $\mathbb{R}^n$!

Beweisskizze:

Man zeigt:

- a) Ein Isomorphismus zwischen endlichdimensionalen Vektorräumen bildet Basen auf Basen ab.
- b) Eine lineare Abbildung zweier gleichdimensionaler Vektorräume, die eine Basis auf eine Basis abbildet, ist ein Isomorphismus.

Kapitel 36

Matrixschreibweise für lineare Abbildungen

36.1 Motivation

- Wir haben gesehen, dass lineare Abbildungen sich durch ihre Wirkung auf die Basisvektoren ausdrücken lassen.
- Mit Hilfe von Matrizen können wir dies kompakt aufschreiben und die Hintereinanderausführung linearer Abbildungen elegant berechnen.

36.2 Struktur linearer Abbildungen zwischen endlichdimensionalen Vektorräumen

Sei K ein Körper (meist $\mathbb{R}$ oder $\mathbb{C}$). Dann ist $f : K^n \rightarrow K^m$

$$\begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \mapsto \begin{pmatrix} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n \\ \vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n \end{pmatrix} \quad (*)$$

eine lineare Abbildung. Für $x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}, y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} \in K^n$ und $\lambda, \mu \in K$ gilt:

$$\begin{aligned} f(\lambda x + \mu y) &= \begin{pmatrix} a_{11}(\lambda x_1 + \mu y_1) & + & \dots & + & a_{1n}(\lambda x_n + \mu y_n) \\ \vdots & & & & \vdots \\ a_{m1}(\lambda x_1 + \mu y_1) & + & \dots & + & a_{mn}(\lambda x_n + \mu y_n) \end{pmatrix} \\ &= \lambda \begin{pmatrix} a_{11}x_1 & + & \dots & + & a_{1n}x_n \\ \vdots & & & & \vdots \\ a_{m1}x_1 & + & \dots & + & a_{mn}x_n \end{pmatrix} \\ &\quad + \mu \begin{pmatrix} a_{11}y_1 & + & \dots & + & a_{1n}y_n \\ \vdots & & & & \vdots \\ a_{m1}y_1 & + & \dots & + & a_{mn}y_n \end{pmatrix} \\ &= \lambda f(x) + \mu f(y). \end{aligned}$$

Umgekehrt kann man sogar zeigen, dass jede lineare Abbildung von K^n nach K^m von dieser Struktur besitzt: Nach 35.9 auf Seite 267 ist eine lineare Abbildung

von K^n nach K^m durch die Bilder der Basisvektoren $\begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \dots, \begin{pmatrix} 0 \\ \vdots \\ 0 \\ 1 \end{pmatrix} \in K^n$ eindeutig bestimmt.

Sind diese Bilder durch $\begin{pmatrix} a_{11} \\ \vdots \\ a_{m1} \end{pmatrix}, \begin{pmatrix} a_{12} \\ \vdots \\ a_{m2} \end{pmatrix}, \dots, \begin{pmatrix} a_{1n} \\ \vdots \\ a_{mn} \end{pmatrix} \in K^m$ gegeben, so ist (*) die entsprechende Abbildung.

Die Struktur (*) motiviert folgende Definition:

36.3 Definition: $((m \times n)$ -Matrix, Zeilenindex, Spaltenindex)

Sei K ein Körper. Das rechteckige Schema

$$A = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix}$$

mit $a_{ij} \in K$ für $i = 1, \dots, m$ und $j = 1, \dots, n$ heißt $(m \times n)$ -Matrix. Die Menge aller $(m \times n)$ -Matrizen über K bezeichnen wir mit $K^{m \times n}$. Die Vektoren

$\begin{pmatrix} a_{11} \\ \vdots \\ a_{m1} \end{pmatrix}, \dots, \begin{pmatrix} a_{1n} \\ \vdots \\ a_{mn} \end{pmatrix}$ heißen Spaltenvektoren von A , und $(a_{11}, \dots, a_{1n}), \dots, (a_{m1}, \dots, a_{mn})$

bilden die Zeilenvektoren. Die Matrix A wird auch als $A = (a_{ij})$ geschrieben. Dabei bezeichnet i den Zeilenindex (bleibt in jeder Zeile konstant) und j ist der Spaltenindex.

36.4 Matrix - Vektor - Produkt

Ist $x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \in K^n$ und $c = \begin{pmatrix} c_1 \\ \vdots \\ c_m \end{pmatrix} \in K^m$, so schreiben wir statt

$$\begin{pmatrix} a_{11}x_1 + \dots + a_{1n}x_n \\ \vdots \\ a_{m1}x_1 + \dots + a_{mn}x_n \end{pmatrix} = \begin{pmatrix} c_1 \\ \vdots \\ c_m \end{pmatrix}$$

jetzt

$$\underbrace{\begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix}}_{\text{Matrix } A} \cdot \underbrace{\begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}}_{\text{Vektor } x} = \underbrace{\begin{pmatrix} c_1 \\ \vdots \\ c_m \end{pmatrix}}_{\text{Vektor } c}.$$

also kurz: $Ax = c$

Eine solche Schreibweise ist sowohl für lineare Abbildungen als auch für lineare Gleichungssysteme (spätere Vorlesung) nützlich. Sie definiert das Produkt zwischen einer Matrix $A = (a_{ij}) \in K^{m \times n}$ und einem Vektor $x \in K^n$ als $A \cdot x = c$ mit $c \in K^m$ und

$$c_i = \sum_{j=1}^n a_{ij}x_j \quad (i = 1, \dots, m)$$

Jede lineare Abbildung $f: K^n \rightarrow K^m$ kann geschrieben werden als $f(x) = Ax$ mit $A \in K^{m \times n}$. Die Spalten von A sind die Bilder der Basisvektoren.

36.5 Beispiele für lineare Abbildung in Matrix-schreibweise

(a) **Streckung:**

Die Matrix $A = \begin{pmatrix} \lambda & 0 \\ 0 & \lambda \end{pmatrix} \in \mathbb{R}^{2 \times 2}$ bildet einen Vektor $\begin{pmatrix} x \\ y \end{pmatrix}$ auf

$$\begin{pmatrix} \lambda & 0 \\ 0 & \lambda \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \lambda \cdot x + 0 \cdot y \\ 0 \cdot x + \lambda \cdot y \end{pmatrix} = \begin{pmatrix} \lambda x \\ \lambda y \end{pmatrix}.$$

(b) **Drehung im $\mathbb{R}^2$:**

$$\begin{pmatrix} 1 \\ 0 \end{pmatrix} \mapsto \begin{pmatrix} \cos \alpha \\ \sin \alpha \end{pmatrix}$$

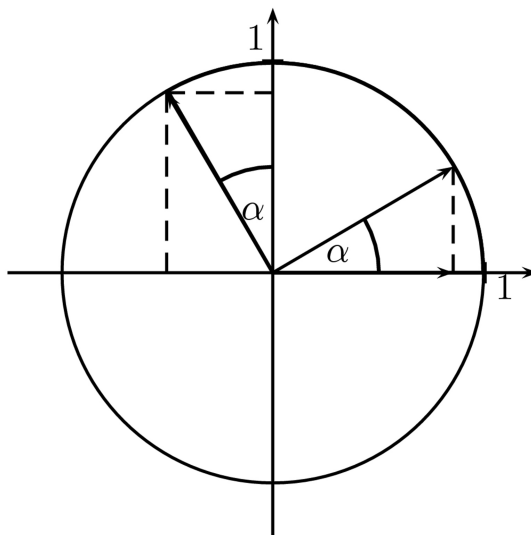
$$\begin{pmatrix} 0 \\ 1 \end{pmatrix} \mapsto \begin{pmatrix} -\sin \alpha \\ \cos \alpha \end{pmatrix}$$

Die Matrix der Drehung um einen Winkel α lautet also

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

Drehungen werden gegen den Uhrzeigersinn durchgeführt.

(c) **Drehung im $\mathbb{R}^3$:**



$$A = \begin{pmatrix} \cos \alpha & 0 & -\sin \alpha \\ 0 & 1 & 0 \\ \sin \alpha & 0 & \cos \alpha \end{pmatrix}$$

beschreibt eine Drehung in der $x_1 - x_3$ - Ebene (entlang der x_2 -Achse passiert nichts)

(d) **Translation:**

$$\begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \mapsto \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} + \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \text{ sind keine linearen Abbildungen!}$$

Welche Rechenoperationen lassen sich mit Matrizen durchführen?

36.6 Definition: (Skalare Multiplikation, Matrixoperationen)

Sei K ein Körper.

- (a) Für $A = (a_{ij}) \in K^{m \times n}$ und $\lambda \in K$ definiert man eine skalare Multiplikation komponentenweise:

$$\lambda A := (\lambda a_{ij}).$$

- (b) Die Addition zweier Matrizen erfolgt ebenfalls komponentenweise:

$$A + B := (a_{ij} + b_{ij}) \quad \forall A = a_{ij} \in K^{m \times n}, B = b_{ij} \in K^{m \times n}.$$

- (c) Bei der Multiplikation von $A \in K^{l \times m}$ mit $B \in K^{m \times n}$ geht man nicht komponentenweise vor:

$$A \cdot B = C \in K^{l \times n} \text{ mit} \\ c_{ij} := \sum_{k=1}^m a_{ik} b_{kj} \quad (\text{„Zeile mal Spalte“})$$

36.7 Beispiele

- (a) $3 \cdot \begin{pmatrix} 5 & 1 \\ 2 & -4 \end{pmatrix} = \begin{pmatrix} 15 & 3 \\ 6 & -12 \end{pmatrix}$
- (b) $\begin{pmatrix} 2 & 3 & 1 \\ 4 & -1 & 7 \end{pmatrix} + \begin{pmatrix} 1 & 5 & -8 \\ 0 & 2 & 9 \end{pmatrix} = \begin{pmatrix} 3 & 8 & -7 \\ 4 & 1 & 16 \end{pmatrix}$
- (c) $\begin{pmatrix} 2 & 3 & 1 \\ 4 & -1 & 7 \end{pmatrix} \cdot \begin{pmatrix} 6 & 4 \\ 1 & 0 \\ 8 & 9 \end{pmatrix} = \begin{pmatrix} 2 \cdot 6 + 3 \cdot 1 + 1 \cdot 8 & 2 \cdot 4 + 3 \cdot 0 + 1 \cdot 9 \\ 4 \cdot 6 - 1 \cdot 1 + 7 \cdot 8 & 4 \cdot 4 - 1 \cdot 0 + 7 \cdot 9 \end{pmatrix} =$
 $\begin{pmatrix} 23 & 17 \\ 79 & 79 \end{pmatrix}$

Welche Rechenregeln folgen aus Def 36.6?

36.8 Satz: (Eigenschaften der Matrixoperationen)

Sei K ein Körper.

- a) $(K^{n \times n}, +, \cdot)$ ist ein Ring mit Eins, d.h. $(K^{n \times n}, +)$ ist eine kommutative Gruppe:

- Assoziativgesetz: $(A + B) + C = A + (B + C)$.
- neutrales Element: $0 = \begin{pmatrix} 0 & \dots & 0 \\ \vdots & & \vdots \\ 0 & \dots & 0 \end{pmatrix} \in K^{n \times n}$ (Nullmatrix).

- *inverses Element* zu A ist $-A = (-1)A$.
- *Kommutativgesetz*: $A + B = B + A$.

$(K^{n \times n}, \cdot)$ ist ein Monoid:

- *Assoziativgesetz*: $(A \cdot B) \cdot C = A \cdot (B \cdot C)$.
- *neutrales Element*:

$$I = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & 1 \end{pmatrix} \in K^{n \times n} \quad \text{Einheitsmatrix.}$$

Es gelten die Distributivgesetze: $A \cdot (B + C) = A \cdot B + A \cdot C$,
 $(A + B) \cdot C = A \cdot C + B \cdot C$.

b) $K^{m \times n}$ ist ein K -Vektorraum:

$(K^{n \times n}, +)$ ist eine kommutative Gruppe:

$$\begin{aligned} \lambda(\mu A) &= (\lambda\mu)A \\ 1 \cdot A &= A \\ \lambda(A + B) &= \lambda A + \lambda B \\ (\lambda + \mu)A &= \lambda A + \mu A \end{aligned}$$

36.9 Bemerkung:

(a) Im Allgemeinen liegt keine Kommutativität bzgl. der Multiplikation vor

$$\begin{pmatrix} 2 & 1 \\ 7 & -4 \end{pmatrix} \cdot \begin{pmatrix} 5 & 3 \\ 4 & 6 \end{pmatrix} = \begin{pmatrix} 2 \cdot 5 + 1 \cdot 4 & 2 \cdot 3 + 1 \cdot 6 \\ 7 \cdot 5 - 4 \cdot 4 & 7 \cdot 3 - 4 \cdot 6 \end{pmatrix} = \begin{pmatrix} 14 & 12 \\ 19 & -3 \end{pmatrix}$$

$$\begin{pmatrix} 5 & 3 \\ 4 & 6 \end{pmatrix} \cdot \begin{pmatrix} 2 & 1 \\ 7 & -4 \end{pmatrix} = \begin{pmatrix} 5 \cdot 2 + 3 \cdot 7 & 5 \cdot 1 - 3 \cdot 4 \\ 4 \cdot 2 + 6 \cdot 7 & 4 \cdot 1 - 6 \cdot 4 \end{pmatrix} = \begin{pmatrix} 31 & -7 \\ 50 & -20 \end{pmatrix}$$

.

(b) Ebenso ist eine $(n \times n)$ -Matrix im Allgemeinen nicht invertierbar bzgl. der Multiplikation.

(c) $(K^{n \times n}, +, \cdot)$ ist nicht nullteilerfrei:

$$\begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} 3 & 7 \\ 0 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}$$

.

(d) Der Vektorraum $K^{n \times m}$ ist isomorph zum Vektorraum $K^{n \cdot m}$.

(e) Spaltenvektoren aus K^m können als $(m \times 1)$ -Matrizen aufgefasst werden. Vektoraddition und skalare Multiplikation sind Spezialfälle der entsprechenden Matrixoperationen.

- (f) Die Hintereinanderausführung linearer Abbildungen entspricht der Multiplikation ihrer Matrizen.

$$K^l \xrightarrow[A \in K^{m \times l}]{g} K^m \xrightarrow[B \in K^{n \times m}]{f} K^n$$

$f \circ g : K^l \rightarrow K^n$ wird durch $B \cdot A \in K^{n \times l}$ repräsentiert.

36.10 Inverse Matrix

Im Allgemeinen hat $A \in K^{n \times n}$ kein multiplikatives Inverses A^{-1} . In vielen Fällen existiert jedoch eine inverse Matrix.

Beispiel: $A = \begin{pmatrix} 5 & 6 \\ 2 & 4 \end{pmatrix}, A^{-1} = \begin{pmatrix} \frac{1}{2} & -\frac{3}{4} \\ -\frac{1}{4} & \frac{5}{8} \end{pmatrix}$, denn

$$\begin{aligned} A \cdot A^{-1} &= \begin{pmatrix} 5 & 6 \\ 2 & 4 \end{pmatrix} \cdot \begin{pmatrix} \frac{1}{2} & -\frac{3}{4} \\ -\frac{1}{4} & \frac{5}{8} \end{pmatrix} \\ &= \begin{pmatrix} 5 & 6 \\ 2 & 4 \end{pmatrix} \cdot \frac{1}{8} \cdot \begin{pmatrix} 4 & -6 \\ -2 & 5 \end{pmatrix} \\ &= \frac{1}{8} \begin{pmatrix} 5 \cdot 4 - 2 \cdot 6 & -5 \cdot 6 + 6 \cdot 5 \\ 2 \cdot 4 - 4 \cdot 2 & -2 \cdot 6 + 4 \cdot 5 \end{pmatrix} \\ &= \frac{1}{8} \begin{pmatrix} 8 & 0 \\ 0 & 8 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}. \end{aligned}$$

Ein Verfahren zur Berechnung der inversen Matrix werden wir in einer späteren Vorlesung kennenlernen.

36.11 Definition: (Invertierbarkeit einer Matrix)

Eine Matrix $A \in K^{n \times n}$ heißt invertierbar (umkehrbar, regulär), falls eine Matrix $A^{-1} \in K^{n \times n}$ existiert mit $A \cdot A^{-1} = I$. Die Menge der invertierbaren Matrizen aus $K^{n \times n}$ wird mit $\text{GL}(n, K)$ bezeichnet.

36.12 Satz: (Gruppeneigenschaft von $\text{GL}(n, K)$)

$\text{GL}(n, K)$ bildet mit der Matrizenmultiplikation eine multiplikative (nichtkommutative) Gruppe.

36.13 Bemerkungen:

- (a) $GL(n, K)$ steht für general linear group.
- (b) Sind $A, B \in GL(n, K)$ so ist $(A \cdot B)^{-1} = B^{-1} \cdot A^{-1}$, denn:
$$A \cdot \underbrace{B \cdot B^{-1}}_I \cdot A^{-1} = A I A^{-1} = A A^{-1} = I.$$
- (c) Ist $n \neq m$, so nennt man $A \in K^{m \times n}$ eine nichtquadratische Matrix. Nicht-quadratische Matrizen sind niemals invertierbar. Allerdings kann man eine so genannte Pseudoinverse angeben. ($\rightarrow$ spätere Vorlesung)
- (d) Die Invertierbarkeit einer Matrix entspricht der Bijektivität ihrer linearen Abbildung. Nach dem Beweis von Satz 35.10 auf Seite 267 ist eine lineare Abbildung zwischen endlichdimensionalen Vektorräumen genau dann bijektiv, wenn Basen auf Basen abgebildet werden.

Folgerung: $A \in K^{n \times n}$ ist genau dann invertierbar, wenn die Spalten von A eine Basis des K^n bilden (d.h. wenn sie linear unabhängig sind).

Kapitel 37

Der Rang einer Matrix

37.1 Motivation

- Interpretiert man eine Matrix als lineare Abbildung, so haben die Spaltenvektoren eine besondere Bedeutung: Sie sind die Bilder der Basisvektoren.
- Wir wollen die lineare Abhängigkeit/Unabhängigkeit dieser Spaltenvektoren genauer untersuchen. Dies liefert u. A. ein wichtiges Kriterium für die Invertierbarkeit von Matrizen.

37.2 Definition: (Spaltenrang)

Unter dem (Spalten-)Rang einer Matrix $A \in K^{m \times n}$ versteht man die maximale Anzahl linear unabhängiger Spaltenvektoren (Schreibweise: $\text{rang } A$).

Welche Aussagen gelten für den Rang?

37.3 Satz: (Aussagen über den Rang einer Matrix)

Ist $f : K^n \rightarrow K^m$ eine lineare Abbildung mit zugehöriger Matrix $A \in K^{m \times n}$ und Spaltenvektoren $a_{*1}, a_{*2}, \dots, a_{*n}$, so gilt:

- (a) $\text{Im}(f) = \text{Span}(a_{*1}, a_{*2}, \dots, a_{*n})$.
- (b) $\dim \text{Im}(f) = \text{rang } A$.
- (c) $\dim \ker(f) = n - \text{rang } A$.

Beweis:

- (a) „ $\supset$ “: Klar, da $a_{*1}, \dots, a_{*n}$ die Bilder der Basisvektoren sind.

$$\text{„}\subset\text{“: } f(x) = Ax = \begin{pmatrix} a_{11}x_1 + \dots + a_{1n}x_n \\ \vdots \\ a_{m1}x_1 + \dots + a_{mn}x_n \end{pmatrix} = x_1 a_{*1} + \dots + x_n a_{*n}.$$

- (b) Folgt direkt aus (a).
(c) Folgt aus Satz 35.5 auf Seite 264.

□

37.4 Beispiele

- (a) Nach 36.13 (d) ist $A \in K^{n \times m}$ genau dann invertierbar, wenn $\text{rang } A = n$ ist.

(b) $A = \begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix}$ hat Rang 3:

$$\begin{aligned} a_{*2}, a_{*3} &\notin \text{Span}(a_{*1}) \\ a_{*3} &\notin \text{Span}(a_{*1}, a_{*2}) \end{aligned}$$

(c) $A = \begin{pmatrix} 1 & 0 & 1 \\ 0 & 1 & 2 \\ 1 & 0 & 2 \end{pmatrix}$ hat Rang 2:

$$\begin{aligned} a_{*1}, a_{*2} &\text{ sind linear unabhängig und} \\ a_{*3} &= a_{*1} + 2 \cdot a_{*2}. \end{aligned}$$

37.5 Definition: (Transponierte Matrix)

Vertauscht man bei einer Matrix $A \in K^{m \times n}$ die Rolle von Zeilen und Spalten, so entsteht die transponierte Matrix $A^T \in K^{n \times m}$.

Der Rang von A^T beschreibt die maximale Zahl der linear unabhängigen Zeilen von A . Man nennt ihn daher auch den Zeilenrang von A .

37.6 Beispiel:

$$A = \begin{pmatrix} 1 & 0 & 1 & 2 \\ 0 & 1 & 2 & 1 \\ 2 & 0 & 2 & 4 \end{pmatrix} \Rightarrow A^T = \begin{pmatrix} 1 & 0 & 2 \\ 0 & 1 & 0 \\ 1 & 2 & 2 \\ 2 & 1 & 4 \end{pmatrix} =: B$$

$\text{rang } A = 2$: a_{*1}, a_{*2} linear unabhängig.

$$a_{*3} = a_{*1} + 2 \cdot a_{*2}$$

$$a_{*4} = 2 \cdot a_{*1} + a_{*2}$$

$\text{rang } B = 2$: b_{*1}, b_{*2} linear unabhängig.

$$b_{*3} = 2 \cdot b_{*1}.$$

In diesem Beispiel stimmen also Spaltenrang und Zeilenrang überein. Ist dies Zufall?

Man kann zeigen:

37.7 Satz: *(Gleichheit von Spaltenrang und Zeilenrang)*

Für $A \in K^{m \times n}$ sind Spalten- und Zeilenrang identisch.

Kapitel 38

Gauss'scher Algorithmus und lineare Gleichungssysteme

38.1 Motivation

- In Abschnitt 36.4 (Matrix-Vektor-Produkte) haben wir gesehen, dass Matrizen zur kompakten Notation linearer Gleichungssysteme benutzt werden können: Die Gleichheit $Ax = b$ (A Matrix, x, b Vektoren) verkörpert die Gleichungen

$$\begin{aligned}a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n &= b_1 \\ &\vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n &= b_m\end{aligned}$$

Lineare Gleichungssysteme treten in einer Vielzahl von Anwendungen auf und müssen gelöst werden.

Vorausgesetzt, A ist eine invertierbare Matrix (vgl. 36.10 - 36.13), dann erhalten wir durch Multiplikation der Gleichung $Ax = b$ mit A^{-1} auf der linken Seite den Vektor x , d.h. inverse Matrizen spielen eine Rolle im Zusammenhang mit der Lösung linearer Gleichungssysteme.

- Aus 37.4 kennen wir den Zusammenhang zwischen Rang und Invertierbarkeit

$$A(n \times n\text{-Matrix}) \text{ ist invertierbar} \Leftrightarrow \text{rang } A = n.$$

Aus diesen Gründen ist es daher von Interesse, ein Verfahren zur Hand zu haben, das

- den Rang einer Matrix ermittelt

- die Inverse (sofern vorhanden) berechnet
- Lineare Gleichungssysteme löst.

Dieses alles leistet der Gauss -Algorithmus.

38.2 Idee

In Beispiel 37.4 (b) $\left(A = \begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix} \right)$ wurde eine obere Dreiecksmatrix betrachtet, die auf der Diagonalen nur von 0 verschiedene Einträge hatte. Allgemein hatte eine $(n \times n)$ -Matrix mit dieser Eigenschaft stets den Rang n .

Die Idee des Gauß-Algorithmus besteht darin, eine beliebige Matrix in eine Dreiecksmatrix umzuwandeln, bzw. eine ähnliche Form, aus der der Rang leicht abzulesen ist, und zwar auf eine Weise, die sicher stellt, dass sich der Rang der Matrix dabei nicht ändert.

38.3 Definition: (elementare Zeilenumformungen)

Die folgenden Operationen auf Matrizen heißen elementare Zeilenumformungen

- (a) Vertauschen zweier Zeilen,
- (b) Addition des λ -fachen der Zeile a_{j*} zur Zeile a_{i*} ($\lambda \in \mathbb{R}$)
- (c) Multiplikation einer Zeile mit einer reellen Zahl $\lambda \neq 0$.

38.4 Satz:

Bei elementaren Zeilenumformungen bleibt der Rang einer Matrix erhalten.

Beweis:

- (a) offensichtlich, da $\text{Span}(a_{1*}, a_{2*}, \dots, a_{m*})$ unverändert bleibt.
- (b) Wir zeigen $\text{Span}(a_{1*}, a_{2*}, \dots, a_{m*}) = \text{Span}(a_{1*}, \dots, a_{i*} + \lambda a_{j*}, \dots, a_{m*})$
„ $\subset$ “: Es sei $v \in \text{Span}(a_{1*}, \dots, a_{i*}, \dots, a_{m*})$, also
$$\begin{aligned} v &= \lambda_1 a_{1*} + \dots + \lambda_i a_{i*} + \dots + \lambda_j a_{j*} + \dots + \lambda_m a_{m*} \\ &= \lambda_1 a_{1*} + \dots + \lambda_i (a_{i*} + \lambda a_{j*}) + \dots + (\lambda_j - \lambda \lambda_i) a_{j*} + \dots + \lambda_m a_{m*} \\ &\in \text{Span}(a_{1*}, \dots, a_{i*} + \lambda a_{j*}, \dots, a_{m*}) \end{aligned}$$
„ $\supset$ “: wird analog gezeigt

(c) analog zu (b).

□

38.5 Gauss'scher Algorithmus

Gegeben sei eine $(m \times n)$ -Matrix $A = (a_{ij})$.

- Betrachte a_{11}
 - Ist $a_{11} \neq 0$, so subtrahiere von jeder Zeile $a_{i*}, i \geq 2$, das $\frac{a_{i1}}{a_{11}}$ -fache der ersten Zeile. Danach ist a_{11} das einzige von Null verschiedene Element der ersten Spalte.
 - Ist $a_{11} = 0$, aber das erste Element a_{i1} einer anderen Zeilen $\neq 0$, so vertausche diese beiden Zeilen (a_{1*} und a_{i*}). Mit der neuen Matrix verfähre wie oben.
 - Sind alle $a_{i1}, i = 1, \dots, m$ gleich 0, so beachte die erste Spalte nicht weiter und verfähre wie oben mit a_{12} (statt a_{11}), sofern auch die 2. Spalte nur Nullen enthält, gehe zur 3. Spalte usw.

Im Ergebnis dieses 1. Schrittes erhalten wir eine Matrix

$$A' = \begin{pmatrix} a'_{11} & a'_{12} & \dots & a'_{1n} \\ 0 & a'_{22} & \dots & a'_{2n} \\ & \vdots & & \vdots \\ 0 & a'_{m2} & \dots & a'_{mn} \end{pmatrix}$$

(evtl. mit zusätzlichen Nullspalten links)

- Nun wird die erste Zeile und Spalte nicht mehr weiter betrachtet und auf die verbleibende Teilmatrix

$$\begin{pmatrix} a'_{22} & \dots & a'_{2n} \\ \vdots & \ddots & \vdots \\ a'_{m2} & \dots & a'_{mn} \end{pmatrix}$$

dasselbe Verfahren angewendet. Damit wird A' umgewandelt in

$$A'' = \begin{pmatrix} a'_{11} & a'_{12} & a'_{13} & \dots & a'_{1n} \\ 0 & a''_{22} & a''_{23} & \dots & a''_{2n} \\ & 0 & a''_{33} & \dots & a''_{3n} \\ & & \vdots & & \vdots \\ 0 & a''_{m3} & \dots & a''_{mn} \end{pmatrix}$$

- Danach betrachtet man wiederum die um eine Zeile und Spalte verkleinerte Teilmatrix usw. bis die gesamte Matrix auf eine Form A^* wie die folgende gebracht ist.

$$A^* = \left(\begin{array}{c|cc|c|cc} a & * & * & & \dots & * \\ \hline 0 & b & * & & \dots & * \\ 0 & 0 & 0 & c & * & * \\ 0 & 0 & 0 & 0 & d & * \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 & 0 & \dots & 0 \end{array} \right) \quad (* \text{ beliebiger Wert})$$

In dieser Matrix heißt der erste von 0 verschiedene Eintrag einer Zeile (a, b, c, d) Leitkoeffizient. Im gesamten Bereich unterhalb und links eines Leitkoeffizienten enthält die Matrix nur Nullen.

$$\left(\begin{array}{ccc|c|ccc} a & * & * & & \dots & * \\ 0 & b & * & & \dots & * \\ \hline 0 & 0 & 0 & c & & * \\ 0 & 0 & 0 & 0 & d & * \\ \vdots & & & \vdots & & * \\ 0 & 0 & 0 & 0 & & * \end{array} \right)$$

Eine solche Matrix heißt Matrix in Zeilen-Stufen-Form.

Formulierung als rekursive Funktion

<pre> Gauss(i, j): falls i = m oder j > n Ende falls a_{ij} = 0: suche a_{kj} ≠ 0, k > i; wenn dies nicht existiert: Gauss(i, j + 1) Ende vertausche Zeile i mit Zeile k für alle k > i: subtrahiere $\frac{a_{kj}}{a_{ij}}$. Zeile i von Zeile k Gauss(i + 1, j + 1) Ende.</pre>	(*)
---	-----

Das Matrixelement a_{ij} das in der Zeile (*) als Nenner auftritt, heißt Pivotelement.

38.6 Beispiel:

$$A = \begin{pmatrix} 0 & 4 & 3 & 1 & 5 \\ 1 & 3 & 2 & 4 & 2 \\ 3 & 1 & 2 & 1 & 0 \\ 2 & 2 & 3 & -2 & 3 \end{pmatrix}$$

Tausche 1. und 2. Zeile

$\Downarrow$

$$= \begin{pmatrix} 1 & 3 & 2 & 4 & 2 \\ 0 & 4 & 3 & 1 & 5 \\ 3 & 1 & 2 & 1 & 0 \\ 2 & 2 & 3 & -2 & 3 \end{pmatrix}$$

Addiere zur 3. Zeile das (-3)-fache der 1. Zeile und
zur 4. Zeile das (-2)-fache der 1. Zeile.

$\Downarrow$

$$A' = \left(\begin{array}{c|cccc} 1 & 3 & 2 & 4 & 2 \\ 0 & 4 & 3 & 1 & 5 \\ 0 & -8 & -4 & -11 & -6 \\ 0 & -4 & -1 & -10 & -1 \end{array} \right)$$

Addiere zur 3. Zeile das 2-fache der 2. Zeile und
zur 4. Zeile das 1-fache der 2. Zeile.

$\Downarrow$

$$A'' = \left(\begin{array}{c|cccc} 1 & 3 & 2 & 4 & 2 \\ 0 & 4 & 3 & 1 & 5 \\ 0 & 0 & 2 & -9 & 4 \\ 0 & 0 & 2 & -9 & 4 \end{array} \right)$$

Addiere zur 4. Zeile das (-1)-fache der 3. Zeile.

$\Downarrow$

$$A''' = \left(\begin{array}{c|cccc} 1 & 3 & 2 & 4 & 2 \\ 0 & 4 & 3 & 1 & 5 \\ 0 & 0 & 2 & -9 & 4 \\ 0 & 0 & 0 & 0 & 0 \end{array} \right) = A^*$$

38.7 Satz:

Der Rang einer Matrix in Zeilen-Stufen-Form ist gleich der Anzahl ihrer Leitkoeffizienten.

Beweis:

Die Anzahl der Leitkoeffizienten sei L .

Jede Zeile, die ein Leitkoeffizienten enthält, ist linear unabhängig von dem System der darunter stehenden Zeilen.

Nimmt man also von unten nach oben die linear unabhängigen Zeilen zur Basis hinzu, so folgt, dass $\dim \text{Span}(a_{1*}, \dots, a_{m*}) \geq L$.
Andererseits ist auch $\dim \text{Span}(a_{1*}, \dots, a_{m*}) \leq L$,

da es nur L Zeilenvektoren $\neq 0$ gibt.

□

38.8 Satz:

Jede elementare Zeilenumformung für $m \times n$ -Matrizen kann als Multiplikation von links mit einer geeigneten invertierbaren $m \times m$ -Matrix D ausgedrückt werden:

(a) Vertauschung Zeile i und j .

$$\begin{pmatrix} 1 & & & & & & & \\ & \ddots & & & & & & \\ & & 1 & & & & & 0 \\ & & & 0 & \dots & 1 & & \\ & & & & 1 & & & \\ & \vdots & & \vdots & \ddots & \vdots & & \vdots \\ & & & & & 1 & & \\ & & & 1 & \dots & 0 & & \\ & & & & & & 1 & \\ & 0 & & & \dots & & & \ddots \\ & & & & & & & & 1 \end{pmatrix} \begin{matrix} \\ \\ \leftarrow i \\ \\ \leftarrow j \\ \\ \\ \\ \end{matrix}$$

$$\begin{matrix} \uparrow & \uparrow \\ i & j \end{matrix}$$

(b) Addition von $\lambda \cdot$ Zeile j zu Zeile i .

$$\begin{pmatrix} 1 & 0 & \dots & \dots & \dots & \dots & 0 \\ 0 & \ddots & & & & & \vdots \\ \vdots & \ddots & 1 & \dots & \boxed{\lambda} & & \\ & & & \ddots & \vdots & & \vdots \\ \vdots & & & & 1 & & \\ \vdots & & & & & \ddots & 0 \\ 0 & \dots & \dots & \dots & \dots & 0 & 1 \end{pmatrix} \begin{matrix} \\ \\ \leftarrow i \\ \\ \leftarrow j \\ \\ \end{matrix}$$

(c) Multiplikation von Zeile i mit $\lambda \neq 0$:

$$\begin{pmatrix} 1 & 0 & \dots & & \dots & \dots & 0 \\ 0 & \ddots & \ddots & & & & \vdots \\ \vdots & \ddots & 1 & & \ddots & & \vdots \\ & & & \lambda & & & \\ \vdots & & \ddots & & 1 & \ddots & \vdots \\ \vdots & & & \ddots & \ddots & \ddots & 0 \\ 0 & \dots & \dots & & \dots & 0 & 1 \end{pmatrix} \leftarrow i$$

(ohne Beweis)

38.9 Umformung einer invertierbaren Matrix zur Einheitsmatrix

- Es sei A eine invertierbare $n \times n$ -Matrix. Wegen $\text{rang } A = n$ hat sie die Zeilen-Stufen-Form

$$\begin{pmatrix} a_{11} & * & \dots & \dots & * \\ 0 & a_{22} & \ddots & & \vdots \\ \vdots & \ddots & \ddots & \ddots & \vdots \\ \vdots & & \ddots & \ddots & * \\ 0 & \dots & \dots & 0 & a_{nn} \end{pmatrix}$$

mit $a_{ii} \neq 0, i = 1, \dots, n$.

- Multiplikation jeder Zeile (a_{i*}) mit $\frac{1}{a_{ii}}$ ergibt eine obere Dreiecksmatrix, deren Diagonalelemente sämtlich gleich 1 sind.
- Mittels weiterer elementarer Zeilenumformungen wird die Matrix in die Einheitsmatrix umgewandelt:
 - Subtrahiere für alle $i < n$ das a_{in} -fache der Zeile n von Zeile i . Danach ist $a_{nn} = 1$ das einzige Element $\neq 0$ in der letzten Spalte
 - Subtrahiere für alle $i < n - 1$ das $a_{i,n-1}$ -fache der Zeile $n - 1$ von Zeile i , usw.

38.10 Satz: (Berechnung der inversen Matrix)

Wird eine invertierbare $n \times n$ -Matrix A durch elementare Zeilenumformungen in die Einheitsmatrix E umgeformt, so erhält man die Inverse A^{-1} , indem man dieselben Umformungen auf die Einheitsmatrix anwendet.

Beweis:

Es seien $D_1, \dots, D_k$ die Matrizen, die gemäß 38.8 die elementaren Zeilenumformungen darstellen, durch die A in E übergeht. Dann ist

$$E = D_k D_{k-1} \dots D_1 A = (D_k D_{k-1} \dots D_1) A$$

Damit ist

$$A^{-1} = D_k D_{k-1} \dots D_1 = D_k D_{k-1} \dots D_1 E.$$

□

38.11 Beispiel

$$\begin{array}{cc} A & E \\ \left(\begin{array}{ccc} 1 & 0 & 2 \\ 2 & -1 & 3 \\ 4 & 1 & 8 \end{array} \right) & \left(\begin{array}{ccc} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{array} \right) \end{array}$$

Addiere zur 2. Zeile das (-2)-fache der 1. Zeile und
zur 3. Zeile das (-4)-fache der 1. Zeile

$$\Downarrow \begin{array}{cc} \left(\begin{array}{ccc} 1 & 0 & 2 \\ 0 & -1 & -1 \\ 0 & 1 & 0 \end{array} \right) & \left(\begin{array}{ccc} 1 & 0 & 0 \\ -2 & 1 & 0 \\ -4 & 0 & 1 \end{array} \right) \end{array}$$

Addiere zur 3. Zeile das 1-fache der 2. Zeile
zur 4. Zeile das 1-fache der 2. Zeile

$$\Downarrow \begin{array}{cc} \left(\begin{array}{ccc} 1 & 0 & 2 \\ 0 & -1 & -1 \\ 0 & 0 & -1 \end{array} \right) & \left(\begin{array}{ccc} 1 & 0 & 0 \\ -2 & 1 & 0 \\ -6 & 1 & 1 \end{array} \right) \end{array}$$

Multipliziere 2. und 3. Zeile mit -1

$$\Downarrow \begin{array}{cc} \left(\begin{array}{ccc} 1 & 0 & 2 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{array} \right) & \left(\begin{array}{ccc} 1 & 0 & 0 \\ 2 & -1 & 0 \\ 6 & -1 & -1 \end{array} \right) \end{array}$$

Addiere zur 1. Zeile das (-2)-fache der 3. Zeile und
zur 1. Zeile das (-1)-fache der 3. Zeile

$$\Downarrow \begin{array}{cc} \left(\begin{array}{ccc} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{array} \right) & \left(\begin{array}{ccc} -11 & 2 & 2 \\ -4 & 0 & 1 \\ 6 & -1 & -1 \end{array} \right) \\ E & A^{-1} \end{array}$$

Probe:

$$\begin{pmatrix} -11 & 2 & 2 \\ -4 & 0 & 1 \\ 6 & -1 & -1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 & 2 \\ 2 & -1 & 3 \\ 4 & 1 & 8 \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Wie kann man den Gauß-Algorithmus zum Lösen linearer Gleichungssysteme verwenden?

38.12 Lineare Gleichungssysteme

Ein lineares Gleichungssystem mit m Gleichungen und n Unbekannten hat die Form $Ax = b$ mit $A \in \mathbb{R}^{m \times n}$, $x \in \mathbb{R}^n$, $b \in \mathbb{R}^m$.

Falls $b \neq 0$, spricht man von einem inhomogenen linearen Gleichungssystem. $Ax = 0$ (d.h. $b = 0$) heißt zugehöriges homogenes Gleichungssystem.

Die Matrix $(A, b) \in \mathbb{R}^{m \times (n+1)}$, d.h. A mit rechts angefügter Spalte b , heißt erweiterte Matrix des Systems.

Interpretiert man A als lineare Abbildung, so sind Lösungen von $Ax = b$ gerade die Vektoren, die durch A und b abgebildet werden.

38.13 Satz: (Lösungsverhalten linearer Gleichungssysteme)

- (a) Die Lösungsmenge des homogenen Systems $Ax = 0$ ist $\ker A$ und daher ein Unterraum von $\mathbb{R}^n$.
- (b) Es sind äquivalent:
 - (i) $Ax = b$ hat mindestens eine Lösung.
 - (ii) $b \in \operatorname{Im} A$.
 - (iii) $\operatorname{rang} A = \operatorname{rang} (A, b)$.
- (c) Ist w eine Lösung von $Ax = b$, so ist die vollständige Lösungsmenge gleich $w + \ker A := \{w + x \mid x \in \ker A\}$

Beweis:

- (a) Klar nach Def. von $\ker A$.
- (b) (i) $\Rightarrow$ (ii): Existiert eine Lösung w , so gilt $Aw = b$, d.h. $b \in \operatorname{Im} A$.
(ii) $\Rightarrow$ (i): Ist $b \in \operatorname{Im} A$, so ist jedes Urbild w von b Lösung von $Ax = b$.
(ii) $\Rightarrow$ (iii): Ist $b \in \operatorname{Im} A$, so ist b Linearkombination der Spalten von A . Damit ist $\operatorname{rang} (A, b) = \operatorname{rang} A$.
(iii) $\Rightarrow$ (ii): Ist $\operatorname{rang} (A, b) = \operatorname{rang} A$, so ist b Linearkombination der Spalten von A . Somit ist $b \in \operatorname{Im} A$.

(c) „ $\subset$ “: Sei $y \in \ker A$ und sei w Lösung von $Ax = b$. Dann gilt:

$$A(w + y) = Aw + Ay = b + 0 = b$$

d.h. $w + y$ ist Lösung von $Ax = b$.

Also liegt $w + \ker A$ in der Lösungsmenge von $Ax = b$.

„ $\supset$ “: Sei v eine weitere Lösung von $Ax = b$. Dann gilt:

$$A(v - w) = Av - Aw = b - b = 0$$

d.h. $v - w \in \ker A$.

$$\Rightarrow v \in w + \ker A$$

Somit liegt die Lösungsmenge von $Ax = b$ in $w + \ker A$.

38.14 Bemerkungen:

- (a) Ist $\text{rang}(A, b) > \text{rang } A$, so hat $Ax = b$ keine Lösung. (folgt aus 38.13).
- (b) Jedes homogene System hat mindestens eine Lösung: 0.
- (c) Zur Lösung des inhomogenen Systems benötigt man
 - die vollständige Lösungsmenge des homogenen Systems,
 - eine spezielle Lösung des inhomogenen Systems.
- (d) Aus (b) und (c) folgt: Das inhomogene System hat genau dann eine eindeutige Lösung, falls $\ker A = 0$ ist.

Um mit Hilfe des Gauß- Algorithmus' das Lösungsverhalten von $Ax = b$ zu studieren, bezeichnen wir mit

$$\text{Lös}(A, b)$$

die Lösungsmenge von $Ax = b$. Wir benötigen

38.15 Lemma: (Invarianz der Lösungsmenge unter Matrixmultiplikation)

Ist $B \in \text{GL}(m, \mathbb{R})$ und $A \in \mathbb{R}^{m \times n}$, $b \in \mathbb{R}^m$, so gilt:

$$\text{Lös}(A, b) = \text{Lös}(BA, Bb)$$

Beweis:

„ $\Rightarrow$ “: Ist $x \in \text{Lös}(A, b)$, gilt $Ax = b$.
 $\Rightarrow BAx = Bb \Rightarrow x \in \text{Lös}(BA, Bb)$.

„ $\Leftarrow$ “: Sei $x \in \text{Lös}(BA, Bb) \Rightarrow BAx = Bb$.
Da $b \in \text{GL}(m, \mathbb{R})$ existiert $B^{-1} \in \text{GL}(m, \mathbb{R})$ nach Satz 36.12
 $\Rightarrow B^{-1}BAx = B^{-1}Bb$
 $\Rightarrow Ax = b \Rightarrow x \in \text{Lös}(A, b)$.

□

Mit diesem Lemma gilt:

38.16 Satz: (*Invarianz der Lösungsmenge unter elementaren Zeilenumformungen*)

Ist die Matrix A' aus A durch elementare Zeilenumformungen entstanden, und der Vektor b' aus b durch die gleichen Zeilenumformungen, so gilt:

$$\text{Lös}(A', b') = \text{Lös}(A, b).$$

Beweisidee:

Elementare Zeilenumformungen entsprechen nach 38.8 der Multiplikation mit invertierbaren Matrizen $D_k, \dots, D_1$.

□

Nun zum eigentlichen Thema

38.17 Gauß-Algorithmus zum Lösen linearer Gleichungssysteme

Zur Lösung von $Ax = b$ geht man in vier Schritten vor:

Schritt 1: Bringe die Matrix (A, b) in Gauß-Jordan-Form (A', b') .

Die Gauß-Jordan-Form ist eine spezielle Zeilen-Stufen-Form, bei der alle Leitkoeffizienten 1 sind und oberhalb der Leitkoeffizienten nur Nullen stehen.

Beispiel:

$$(A', b') = \left(\begin{array}{cccccc|c} 1 & 0 & 3 & 0 & 0 & 8 & 2 \\ 0 & 1 & 2 & 0 & 0 & 1 & 4 \\ 0 & 0 & 0 & 1 & 0 & 5 & 6 \\ 0 & 0 & 0 & 0 & 1 & 4 & 0 \end{array} \right) \text{ ist in Gauß - Jordan Form.}$$

$x_1 \quad x_2 \quad x_3 \quad x_4 \quad x_5 \quad x_6$

$\text{rang } A' = 4$ und $\text{rang } (A', b') = 4$

$Ax = b$ hat (mind.) eine Lösung.

$\Rightarrow$ Es existieren Lösungen, wir müssen weiter machen.

Schritt 2: Finde die Lösungsmenge U des homogenen Gleichungssystem $A'x = 0$.
Wähle hierzu die Unbekannten, die zu den Spalten ohne Leitkoeffizienten gehören, als freie Parameter.

Beispiel:

Im obigen Beispiel setzen wir $x_3 := \lambda, x_6 := \mu$.

Dann hat das homogene System die Lösungsmenge

$$\begin{aligned} x_1 &= -3\lambda - 8\mu \\ x_2 &= -2\lambda - \mu \\ x_3 &= \lambda \\ x_4 &= -5\mu \\ x_5 &= -4\mu \\ x_6 &= \mu \end{aligned}$$

Als Vektorgleichung:

$$\begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \\ x_6 \end{pmatrix} = \begin{pmatrix} -3\lambda - 8\mu \\ -2\lambda - \mu \\ \lambda \\ -5\mu \\ -4\mu \\ \mu \end{pmatrix} = \lambda \begin{pmatrix} -3 \\ -2 \\ 1 \\ 0 \\ 0 \\ 0 \end{pmatrix} + \mu \begin{pmatrix} -8 \\ -1 \\ 0 \\ -5 \\ -4 \\ 1 \end{pmatrix}$$

$$U = \text{Span} \left(\begin{pmatrix} -3 \\ -2 \\ 1 \\ 0 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} -8 \\ -1 \\ 0 \\ -5 \\ -4 \\ 1 \end{pmatrix} \right)$$

Schritt 3: Suche eine spezielle Lösung w des inhomogenen Systems $A'x = b'$. Setze hierzu die Unbekannte, die zu den Spalten ohne Leitkoeffizienten gehören, gleich 0. (möglich, da dies freie Parameter sind).

Beispiel: Mit $x_3 = 0, x_6 = 0$ erlaubt die Gauß-Jordan-Form im obigen Beispiel die direkte Bestimmung von x_1, x_2, x_4, x_5 :

$$x_1 = 2, x_2 = 4, x_4 = 6, x_5 = 0.$$

$$\Rightarrow w = \begin{pmatrix} 2 \\ 4 \\ 0 \\ 6 \\ 0 \\ 0 \end{pmatrix}$$

ist spezielle Lösung des inhomogenen Systems.

Schritt 4: Die Lösungsmenge von $Ax = b$ ist dann $w + U$.

Im Beispiel:

$$\text{Lös}(A, b) = \left\{ \begin{pmatrix} 2 \\ 4 \\ 0 \\ 6 \\ 0 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} -3 \\ -2 \\ 1 \\ 0 \\ 0 \\ 0 \end{pmatrix} + \mu \begin{pmatrix} -8 \\ -1 \\ 0 \\ -5 \\ -4 \\ 1 \end{pmatrix} \mid \lambda, \mu \in \mathbb{R} \right\}$$

38.18 Geometrische Interpretation linearer Gleichungssysteme

Sei $A \in \mathbb{R}^{m \times n}, x \in \mathbb{R}^n, b \in \mathbb{R}^m$. Wir wissen: $\text{Lös}(A, 0) = \ker A$ ist Unterraum des $\mathbb{R}^n$.

(Jedoch ist $\text{Lös}(A, b)$ i.A. kein Unterraum)

Für die Dimension von $\text{Lös}(A, 0)$ gilt:

$$\begin{aligned} \underbrace{\dim \ker A}_{\dim \text{Lös}(A, 0)} + \underbrace{\dim \text{Im } A}_{\text{rang } A} &= \underbrace{\dim \mathbb{R}^n}_n \\ \Rightarrow \dim \text{Lös}(A, 0) &= n - \text{rang } A. \end{aligned}$$

Beispiel:

$$A = \begin{pmatrix} 1 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \in \mathbb{R}^{2 \times 3}, b = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$$

$$n = 3, \text{rang } A = 2$$

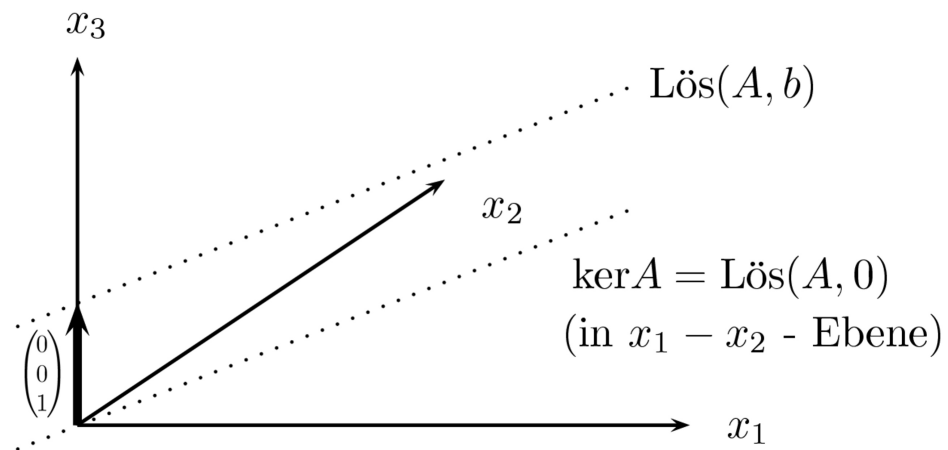
$$\Rightarrow \dim \text{Lös}(A, 0) = 3 - 2 = 1.$$

$$\text{Lös}(A, 0) = \left\{ \begin{pmatrix} \lambda \\ \lambda \\ 0 \end{pmatrix} \mid \lambda \in \mathbb{R} \right\} = \ker A \quad \text{Ursprungsgerade}$$

$$\text{spezielle Lösung des inhomogen Systems: } w = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$$

allgemeine Lösung des inhomogen Systems:

$$\text{Lös}(A, b) = \left\{ \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} + \lambda \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} \mid \lambda \in \mathbb{R} \right\} \quad \text{Gerade.}$$



Iterative Verfahren für lineare Gleichungssysteme

- Viele praktische Probleme führen auf sehr große lineare Gleichungssysteme, bei denen die Systemmatrix dünn besetzt ist, d.h. nur wenige Einträge von 0 verschieden sind.

[illegible]

fusionsfiltern in der Bildverarbeitung.

- Beispiel:

- Direkte Verfahren wie der Gauß-Algorithmus können die Nullen auffüllen und so zu einem enormen Speicherplatzbedarf führen.

Zudem ist der Rechenaufwand oft zu hoch: $O(n^3)$ Operationen für ein $(n \times n)$ -Gleichungssystem.

- Daher verwendet man oft iterative Näherungsverfahren, die kaum zusätzlichen Speicherplatz benötigen und nach wenigen Schritten eine brauchbare Approximation liefern.

39.2 Grundstruktur klassischer iterativer Verfahren

Geg.: $A \in \mathbb{R}^{n \times n}, b \in \mathbb{R}^n$

Ges.: $x \in \mathbb{R}^n$ mit $Ax = b$

Falls $A = S - T$ mit einer einfach zu invertierenden Matrix S (z.B. Diagonalmatrix, Dreiecksmatrix), so kann man $Ax = b$ umformen in

$$Sx = Tx + b$$

und mit einem Startwert $x^{(0)}$ die Fixpunktiteration

$$Sx^{(k+1)} = Tx^{(k)} + b \quad (k = 0, 1, 2, \dots)$$

anwenden.

Wir wollen nun drei verschiedene Aufspaltungen $A = S - T$ untersuchen. Sei hierzu $A = D - L - R$ eine Aufspaltung in eine Diagonalmatrix D , eine strikte untere Dreiecksmatrix L und eine strikte obere Dreiecksmatrix R .

$$\underbrace{\begin{bmatrix} * \\ \\ \\ \end{bmatrix}}_A = \underbrace{\begin{bmatrix} * & & \mathbf{0} \\ & \ddots & \\ \mathbf{0} & & * \end{bmatrix}}_D - \underbrace{\begin{bmatrix} 0 & \dots & 0 \\ & \ddots & \vdots \\ * & & 0 \end{bmatrix}}_L + \underbrace{\begin{bmatrix} 0 & & * \\ \vdots & \ddots & \\ 0 & \dots & 0 \end{bmatrix}}_R -$$

39.3 Das Jacobi-Verfahren (Gesamtschrittverfahren)

Hierzu wählt man $S := D$ und $T := L + R$.

In jeder Iteration wird also nur das Diagonalsystem

$$Dx^{(k+1)} = (L + R)x^{(k)} + b$$

nach $x^{(k+1)}$ aufgelöst, d.h. nur durch die Diagonalelemente dividiert:

$$x^{(k+1)} = D^{-1} \left((L + R)x^{(k)} + b \right)$$

explizit:

$$x_i^{(k+1)} = \frac{1}{a_{ii}} \left(- \sum_{\substack{j=1 \\ j \neq i}}^n a_{ij} x_j^{(k)} + b_i \right) \quad (i = 1, \dots, n)$$

39.4 Beispiel

$$A = \begin{pmatrix} 2 & -1 \\ -1 & 2 \end{pmatrix}, b = \begin{pmatrix} 3 \\ 4 \end{pmatrix}$$

führt auf das Diagonalsystem

$$\left. \begin{aligned} 2x_1^{(k+1)} &= x_2^{(k)} + 3 \\ 2x_2^{(k+1)} &= x_1^{(k)} + 4 \end{aligned} \right\} \Rightarrow \begin{aligned} x_1^{(k+1)} &= \frac{1}{2} (x_2^{(k)} + 3) \\ x_2^{(k+1)} &= \frac{1}{2} (x_1^{(k)} + 4) \end{aligned}.$$

Bei vierstelliger Genauigkeit und Startvektor $x^{(0)} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$ erhält man

k	$x_1^{(k)}$	$x_2^{(k)}$
0	0	0
1	1,5	2
2	2,5	2,75
3	2,875	3,25
4	3,125	3,438
5	3,219	3,568
6	3,282	3,610
7	3,305	3,641
8	3,321	3,653
9	3,327	3,661
10	3,331	3,664
11	3,332	3,666
12	3,333	3,666

Exakte Lösung: $x_1 = 3\frac{1}{3}, x_2 = 3\frac{2}{3}$.

Das Jacobi-Verfahren ist gut geeignet für Parallelrechner, da $x_i^{(k+1)}$ nicht von $x_j^{(k+1)}$ abhängt.

39.5 Das Gauß-Seidel-Verfahren (Einzelschrittverfahren)

$$S := D - L, T := R$$

In jeder Iteration wird also das Dreieckssystem

$$(D - L)x^{(k+1)} = Rx^{(k)} + b$$

durch einfache Vorwärtssubstitution gelöst:

$$x_i^{(k+1)} = \frac{1}{a_{ii}} \left(- \sum_{j=1}^{i-1} a_{ij} x_j^{(k+1)} - \sum_{j=i+1}^n a_{ij} x_j^{(k)} + b_i \right) \quad (i = 1, \dots, n)$$

d.h. neue Werte werden weiter verwendet, sobald sie berechnet wurden.

39.6 Beispiel:

$$A = \begin{pmatrix} 2 & -1 \\ -1 & 2 \end{pmatrix}, b = \begin{pmatrix} 3 \\ 4 \end{pmatrix}, x = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$

$$\begin{aligned} x_1^{(k+1)} &= \frac{1}{2} (x_2^{(k)} + 3) \\ x_2^{(k+1)} &= \frac{1}{2} (x_1^{(k+1)} + 4) \end{aligned}$$

k	$x_1^{(k)}$	$x_2^{(k)}$
0	0	0
1	1,5	2,75
2	2,875	3,348
3	3,174	3,587
4	3,294	3,647
5	3,324	3,662
6	3,331	3,666

Konvergiert etwa doppelt so schnell wie Jacobi-Verfahren.

39.7 Das SOR-Verfahren (SOR = successive over-relaxation)

Beschleunigung des Gauß-Seidel-Verfahrens durch Extrapolation

$$x^{(k+1)} = \tilde{x}^{(k+1)} + \omega (\tilde{x}^{(k+1)} - x^{(k)})$$

mit $\omega \in (1, 2)$, wobei $\tilde{x}^{(k+1)}$ die Gauß-Seidel-Iteration zu $x^{(k)}$ ist.

Für große Gleichungssysteme kann damit die Iterationszahl (um eine bestimmte Genauigkeit zu erreichen) für ein geeignetes ω um eine Zehnerpotenz gesenkt werden.

39.8 Konvergenzresultate

- (a) Die Jacobi-, Gauß-Seidel- und SOR-Verfahren können als Fixpunktiterationen aufgefasst werden.
Nach 26.6 liegt Konvergenz vor, wenn die Abbildung kontrahierend ist. Dies ist nicht in allen Fällen erfüllt.
- (b) Für wichtige, praxisrelevante Fälle existieren jedoch Konvergenzaussagen, z.B.: Ist die Systemmatrix $A = (A_{ij}) \in \mathbb{R}^{n \times n}$ streng diagonaldominant (d.h. $|a_{ii}| > \sum_{\substack{j=1 \\ j \neq i}}^n |a_{ij}| \quad \forall i$), so konvergieren das Jacobi- und das Gauß-Seidel-Verfahren.

39.9 Effizienz

- (a) Die vorgestellten Verfahren sind die einfachsten, allerdings nicht die effizientesten iterativen Verfahren für lineare Gleichungssysteme.
- (b) Effizientere, aber kompliziertere Verfahren :
- ADI - Verfahren (ADI = alternating directions implicit)
(mehrdimensionale Probleme werden in eindimensionale Probleme umgewandelt)
 - PCG - Verfahren (PCG = preconditioned conjugate gradients)
(suchen in Unterräumen)
 - Mehrgitterverfahren (multigrid)
(erst grobe Berechnung, dann feiner)

Siehe numerische Spezialliteratur.

- (c) Hocheffiziente Ansätze wie die Mehrgitterverfahren verwenden oft das Gauß-Seidel-Verfahren als Grundbaustein. Mit ihnen ist es z.T. möglich, lineare Gleichungssysteme in optimaler Komplexität (d.h. in $O(n)$) zu lösen.

Kapitel 40

Determinanten

40.1 Motivation

- Gibt es eine möglichst „aussagekräftige“ Abbildung, die eine Matrix $A \in K^{n \times n}$ auf eine Zahl in K reduziert.
- Die Determinante ist die „sinnvollste“ Art, eine solche Abbildung zu definieren.
- Sie kann axiomatisch fundiert werden und liefert nützliche Aussagen:
 - über die Invertierbarkeit einer Matrix
 - über die Lösung des linearen Gleichungssystems $Ax = b$
 - über das Volumen eines Parallelepipeds

Welche Forderung soll die Determinante erfüllen?

40.2 Definition: (*Determinantenfunktion, Determinante*)

Sei K ein Körper: Eine Abbildung: $\det : K^{n \times n} \rightarrow K$ heißt Determinantenfunktion (Determinante), falls gilt:

(a) $\det$ ist linear in jeder Zeile:

$$\det \begin{pmatrix} z_1 \\ \vdots \\ z_{i-1} \\ \lambda z_i + \mu z'_i \\ z_{i+1} \\ \vdots \\ z_n \end{pmatrix} = \lambda \det \begin{pmatrix} z_1 \\ \vdots \\ z_{i-1} \\ z_i \\ z_{i+1} \\ \vdots \\ z_n \end{pmatrix} + \mu \det \begin{pmatrix} z_1 \\ \vdots \\ z_{i-1} \\ z'_i \\ z_{i+1} \\ \vdots \\ z_n \end{pmatrix}$$

für $i = 1, \dots, n$ und $\lambda, \mu \in K$.

(Man sagt auch die Determinante ist eine Multilinearform)

(b) Ist $\text{rang } A < n$, so ist $\det A = 0$.

(c) $\det I = 1$ für die Einheitsmatrix $I \in K^{n \times n}$.

Bemerkung:

Die Determinante ist nur für quadratische Matrizen definiert!

Gibt es sehr viele Determinantenfunktionen?

Nein! In einem aufwändigen Beweis (siehe z.B. Beutelspacher: Lineare Algebra) kann man zeigen:

40.3 Satz: (Eindeutigkeit der Determinante)

Zu jedem $n \in \mathbb{N}$ existiert genau eine Determinantenfunktion.

Mit anderen Worten: Die Forderungen (a) - (c) aus 40.2 liefern eine axiomatische Fundierung des Determinantenbegriffs.

Wie berechnet man Determinanten? Hierzu betrachten wir zunächst nur 2×2 -Matrizen.

40.4 Satz: (Determinanten einer 2×2 -Matrix)

Die Determinante einer Matrix $\begin{pmatrix} a & b \\ c & d \end{pmatrix} \in K^{2 \times 2}$ ist gegeben durch

$$\det \begin{pmatrix} a & b \\ c & d \end{pmatrix} := \begin{vmatrix} a & b \\ c & d \end{vmatrix} := a \cdot d - b \cdot c.$$

Beweis:

Wir zeigen, dass diese Abbildung Def. 40.2 erfüllt.

(a)

$$\begin{aligned} \det \begin{pmatrix} \lambda a_1 + \mu a_2 & \lambda b_1 + \mu b_2 \\ c & d \end{pmatrix} &= (\lambda a_1 + \mu a_2) \cdot d - (\lambda b_1 + \mu b_2) \cdot c \\ &= \lambda a_1 d + \mu a_2 d - \lambda b_1 c - \mu b_2 c \\ &= \lambda(a_1 d - b_1 c) + \mu(a_2 d - b_2 c) \\ &= \lambda \begin{pmatrix} a_1 & b_1 \\ c & d \end{pmatrix} + \mu \begin{pmatrix} a_2 & b_2 \\ c & d \end{pmatrix} \end{aligned}$$

Die Linearität in der 2. Zeile zeigt man analog.

- (b) Wenn die Matrix nur aus Nullen besteht, ist ihrer Determinante Null. Hat die Matrix rang 1, so ist $\begin{pmatrix} a & b \\ c & d \end{pmatrix} = \lambda \begin{pmatrix} b \\ d \end{pmatrix}$, d.h.

$$\begin{aligned} \det \begin{pmatrix} a & b \\ c & d \end{pmatrix} &= \det \begin{pmatrix} \lambda b & b \\ \lambda d & d \end{pmatrix} \\ &= \lambda b d - \lambda b d = 0. \end{aligned}$$

(c) $\det \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = 1 \cdot 1 - 0 \cdot 0 = 1.$

□

Bemerkung: Für eine 1×1 -Matrix $a \in K^{1 \times 1}$ ist $\det(a) = a$.

Determinanten von $(n \times n)$ -Matrizen lassen sich rekursiv auf 2×2 -Determinanten zurückführen. Hierzu benötigen wir:

40.5 Definition: (Unterdeterminante)

Sei $A \in K^{n \times n}$. Die aus einer $n \times n$ -Determinanten $D = \det A$ durch Streichung der i -ten Zeile und j -ten Spalte entstehende $(n-1) \times (n-1)$ -Determinante D_{ij} nennen wir Unterdeterminante von D .

Der Ausdruck $A_{ij} := (-1)^{i+j} D_{ij}$ heißt algebraisches Komplement des Elementes a_{ij} in der Determinante D .

40.6 Beispiel

$$\begin{aligned} A &= \begin{pmatrix} 3 & 9 & 1 \\ 2 & 5 & 4 \\ -2 & 8 & 7 \end{pmatrix} \\ D_{23} &= \begin{vmatrix} 3 & 9 \\ -2 & 8 \end{vmatrix} = 3 \cdot 8 - (-2) \cdot 9 = 42 \\ A_{23} &= (-1)^{2+3} \cdot D_{23} = -42 \end{aligned}$$

Kommen wir nun zur rekursiven Berechnung einer $n \times n$ -Determinante.

40.7 Satz: (Laplace'scher Entwicklungssatz)

Man kann eine $n \times n$ -Determinante berechnen, indem man die Elemente einer Zeile (oder Spalte) mit ihrem algebraischen Komplemente multipliziert und die Produkte addiert.

40.8 Beispiel:

- (a) Entwicklung einer 3×3 -Determinante nach der 2. Zeile:

$$\begin{vmatrix} 3 & 9 & 1 \\ 2 & 5 & 4 \\ -2 & 8 & 7 \end{vmatrix} = -2 \begin{vmatrix} 9 & 1 \\ 8 & 7 \end{vmatrix} + 5 \begin{vmatrix} 3 & 1 \\ -2 & 7 \end{vmatrix} - 4 \begin{vmatrix} 3 & 9 \\ -2 & 8 \end{vmatrix} \\ = -2(63 - 8) + 5(21 + 2) - 4(24 + 18) = \dots = -163.$$

- (b) Entwicklung nach der 3. Spalte:

$$\begin{vmatrix} 3 & 9 & 1 \\ 2 & 5 & 4 \\ -2 & 8 & 7 \end{vmatrix} = 1 \begin{vmatrix} 2 & 5 \\ -2 & 8 \end{vmatrix} - 4 \begin{vmatrix} 3 & 9 \\ 2 & 8 \end{vmatrix} + 7 \begin{vmatrix} 3 & 9 \\ -2 & 5 \end{vmatrix} \\ = 1(16 + 10) - 4(24 + 18) + 7(15 - 18) = \dots = -163.$$

Wie rechnet man mit Determinanten?

40.9 Rechenregel für Determinanten

- (a) Transponieren verändert den Wert einer Determinante nicht:

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

(folgt aus dem Laplace'schen Entwicklungssatz, indem man die Entwicklung nach Zeilen und Spalten vertauscht).

- (b) Aus Def. 40.2 (b) folgt: Sind Spaltenvektoren oder Zeilenvektoren linear abhängig, so ist die Determinante Null.

Beispiel: $\begin{vmatrix} 1 & 2 & 3 \\ 3 & 3 & 3 \\ 6 & 6 & 6 \end{vmatrix} = 0$

- (c) Additiert man zu einer Zeile (oder Spalte) das Vielfache einer anderen Zeile (Spalte), so bleibt die Determinante gleich.

Beispiel: $\begin{vmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{vmatrix} = \begin{vmatrix} 1 & 2 & 3 \\ 3 & 3 & 3 \\ 6 & 6 & 6 \end{vmatrix} \stackrel{(b)}{=} 0$

(1. Zeile von 2. und 3. abgezogen)

- (d) Vertauscht man zwei Zeilen oder Spalten, so ändert die Determinante ihr Vorzeichen.

- (e) Die Determinante von Dreiecksmatrizen ist das Produkt der Diagonalelemente

(folgt durch rekursives Anwenden des Laplace'schen Entwicklungssatzes):

$$\begin{vmatrix} 3 & 0 & 9 \\ 0 & -7 & 4 \\ 0 & 0 & 2 \end{vmatrix} = 3 \cdot (-7) \cdot 2 = -42$$

Man kann also mit dem Gauß-Algorithmus die Matrix auf Dreiecksgestalt bringen (unter Beachtung von (d)), um dann ihre Determinante bequem zu berechnen. Für große n ist dies wesentlich effizienter als der Laplace'sche Entwicklungssatz.

(f) Für $A, B \in K^{n \times n}$ gilt: $\boxed{\det(A \cdot B) = \det A \cdot \det B}$

(g) Folgerung für eine invertierbare Matrix A gilt:

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

$$\Rightarrow \boxed{\det(A^{-1}) = \frac{1}{\det A}}$$

(h) Vorsicht: Für $A \in K^{n \times n}$, $\lambda \in K$ gilt:

$$\boxed{\det(\lambda A) = \lambda^n \det A}$$

(und nicht etwas $\det(\lambda A) = \lambda \det A$, denn $\det$ ist linear in jeder Zeile).

Wozu sind Determinanten nützlich?

40.10 Bedeutung der Determinanten

(a) Mit ihnen kann man testen, ob eine Matrix $A \in K^{n \times n}$ invertierbar ist:

$$\boxed{A \in K^{n \times n} \text{ ist invertierbar} \Leftrightarrow \det A \neq 0}$$

(b) Man kann mit ihnen lineare Gleichungssysteme lösen.
(für numerische Rechnungen ist dies jedoch ineffizient):

Cramer'sche Regel: Ist $A = (a_{*1}, \dots, a_{*n}) \in K^{n \times n}$ invertierbar und $b \in K^n$, so läßt sich die Lösung des linearen Gleichungssystems $Ax = b$ angeben durch

$$x_k = \frac{\det(a_{*1}, \dots, a_{*k-1}, b, a_{*k+1}, \dots, a_{*n})}{\det A} \quad (k = 1, \dots, n)$$

Beispiel:

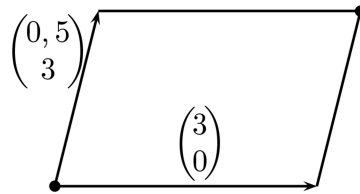
$$\begin{pmatrix} 2 & 5 \\ 1 & 4 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} -2 \\ 6 \end{pmatrix}$$

$$x_1 = \frac{\begin{vmatrix} -2 & 5 \\ 6 & 4 \end{vmatrix}}{\begin{vmatrix} 2 & 5 \\ 1 & 4 \end{vmatrix}} = \frac{-8 - 30}{8 - 5} = \frac{-38}{3}$$

$$x_2 = \frac{\begin{vmatrix} 2 & -2 \\ 1 & 6 \end{vmatrix}}{\begin{vmatrix} 2 & 5 \\ 1 & 4 \end{vmatrix}} = \frac{12 + 2}{3} = \frac{14}{3}$$

- (c) $|\det A|$ ist das Volumen des durch die Spaltenvektoren von A aufgespannten Parallelepipeds.

Beispiel: Parallelogrammfläche:



$$\left| \det \begin{pmatrix} 3 & 0,5 \\ 0 & 2 \end{pmatrix} \right| = |3 \cdot 2 - 0 \cdot 0,5| = |6|.$$

Literaturverzeichnis

- [FOR01] Otto Forster - Analysis 1, 6., verbesserte Auflage, Vieweg Verlag, 2001
- [HAR03] Peter Hartmann - Mathematik für Informatiker, Vieweg Verlag, 2003
- [BRI01] Manfred Brill - Mathematik für Informatiker, Hanser Verlag, 2001
- [DOR88] Willibald Dörfler, Werner Peschek - Einführung in die Mathematik für Informatiker, Hanser Verlag, 1988
- [BEU01] Albrecht Beutelspacher - Lineare Algebra, Vieweg Verlag, 2001
- [HEI00] Heiner Heil, Joachim Kläsner - Analysis (Skript für Leistungskurs), Softfrutti Verlag, 2000
- [SCH??] Wolfgang Scholl, Rainer Drews - Handbuch Mathematik, Orbis Verlag / Falken Verlag, ????