

# VORLESUNGSMODUL MATHEMATIK 1

## - VORLMOD MATHE1 -

MATTHIAS ANSORG

**ZUSAMMENFASSUNG.** Studentische Mitschrift zur Vorlesung »Mathematik 1« bei Prof. Dr. Lutz Eichner (Wintersemester 2002/2003) im Studiengang Informatik an der FH Gießen-Friedberg. Zu dieser Vorlesung Mathematik 1 benötigt man theoretisch keine mathematischen Vorkenntnisse, doch sind diese natürlich sehr hilfreich. Die Übung Mi 17:30h findet nicht statt. Alle anderen Übungen finden statt. Wenn die Übungsgruppe Fr 08:30h zu schlecht besucht ist, wird sie abgeschafft. Die Übungen finden wöchentlich statt. Man sollte an einer Übungsgruppe teilnehmen und bei dieser bleiben. Die Übungsaufgaben stehen auch auf der Homepage von Prof. Dr. Eichner; Benutzername: student, Passwort: student.

- **Bezugsquelle:** freier Download unter <http://matthias.ansorgs.de/InformatikDiplom/Modul.Mathe1.Eichner/Mathe1.pdf>.
- **Lizenz:** Diese studentische Mitschrift ist public domain, darf also ohne Einschränkungen oder Quellenangabe für jeden beliebigen Zweck benutzt werden, kommerziell und nichtkommerziell; jedoch enthält sie keinerlei Garantien für Richtigkeit oder Eignung oder sonst irgendetwas, weder explizit noch implizit. Das Risiko der Nutzung dieser studentischen Mitschrift liegt allein beim Nutzer selbst. Einschränkend sind außerdem die Urheberrechte der verwendeten Quellen zu beachten.
- **Korrekturen:** Fehler zur Verbesserung in zukünftigen Versionen, sonstige Verbesserungsvorschläge und Wünsche bitte dem Autor per e-mail mitteilen: Matthias Ansorg, ansis@gmx.de.
- **Format:** Das vorliegende Script wurde mit dem Programm L<sup>A</sup>T<sub>E</sub>X (graphisches Frontend zu L<sup>A</sup>T<sub>E</sub>X) unter Linux erstellt und als pdf-Datei exportiert.
- **Dozent:** Prof. Dr. Lutz Eichner.
- **Verwendete Quellen:** [1], [2], [11]. Für hieraus übernommene Teile liegt eine Erlaubnis vor.
- **Klausur:** Die Aufgaben der Klausur sind Variationen der Aufgaben in den Übungen. In der Klausur dürfen alle Papier-Unterlagen als Hilfsmittel benutzt werden, aber *kein* Taschenrechner. Auf die Hausübungen gibt es Bonuspunkte; man kann damit die Klausurnote verbessern, für das Bestehen der Klausur ist die Abgabe der Hausübungen jedoch irrelevant. Wichtig: Kopf der Hausübungen korrekt ausfüllen. Erhält man alle mit den Hausübungen erreichbaren 4 Punkte, so wird die Klausurnote dadurch etwa eine Note besser. Noten und Punkte in der Klausur in etwa:
  - Note 1: 32 bis 29 Punkte
  - Note 2: 28 bis 25 Punkte
  - Note 3: 24 bis 21 Punkte
  - Note 4: 20 bis 12 Punkte
  - Note 5: 11 bis 0 Punkte

Es wird eine Klausur vor und eine nach den Ferien geschrieben. Um die Klausur vor den Ferien zu schreiben, muss man mindestens 50% der mit den Hausübungen erreichbaren Punkte erhalten. Die Übungen müssen mit der eigenen Handschrift angefertigt werden. Wenn die Korrektoren den Eindruck haben, dass abgeschrieben wurde, schreiben sie »0 Punkte, abgeschrieben«. Man kann seine Punkte dann trotzdem durch eine Vorführung der Aufgaben bei Prof. Eichner erhalten.

## INHALTSVERZEICHNIS

|                                                          |    |
|----------------------------------------------------------|----|
| Abbildungsverzeichnis                                    | 4  |
| Tabellenverzeichnis                                      | 5  |
| 1. Einige Bezeichnungen                                  | 5  |
| 2. Grundbegriffe der Mengenlehre                         | 5  |
| 2.1. Möglichkeiten, eine Menge anzugeben                 | 6  |
| 2.2. Möglichkeiten, eine Menge darzustellen              | 7  |
| 2.3. Relationen von Mengen (Beziehungen zwischen Mengen) | 7  |
| 2.4. Verknüpfungen von Mengen                            | 9  |
| 2.5. Supersatz                                           | 10 |
| 3. Aussagenlogik                                         | 12 |
| 3.1. Relationen                                          | 12 |
| 3.2. Grundgesetze der Aussagenlogik                      | 14 |
| 4. Beweismethoden                                        | 14 |

|                                                                              |    |
|------------------------------------------------------------------------------|----|
| 4.1. Direkter Beweis einer Behauptung $B$                                    | 15 |
| 4.2. Indirekter Beweis einer Behauptung $B$                                  | 15 |
| 4.3. Beispiele zu Beweismethoden                                             | 15 |
| 5. Abbildungen                                                               | 16 |
| 6. Permutation                                                               | 17 |
| 7. Zahlen                                                                    | 17 |
| 7.1. Grundgesetze der reellen Zahlen                                         | 17 |
| 7.2. Warum gilt $\mathbb{Q} \subset \mathbb{R}$ ?                            | 17 |
| 7.3. Irrationale Zahlen                                                      | 18 |
| 7.4. Algebraische und transzendente Zahlen                                   | 19 |
| 8. Mächtigkeiten von Zahlenmengen                                            | 19 |
| 9. Vollständige Induktion                                                    | 20 |
| 10. Binomialkoeffizienten                                                    | 23 |
| 11. Winkelfunktionen                                                         | 25 |
| 11.1. Winkel                                                                 | 25 |
| 11.2. Einführung                                                             | 26 |
| 11.3. Definition der Winkelfunktionen für beliebige Winkel                   | 28 |
| 11.4. Bogenmaß                                                               | 29 |
| 11.5. Graphen der Winkelfunktionen                                           | 29 |
| 11.6. Additionstheoreme                                                      | 30 |
| 12. Umkehrfunktion                                                           | 30 |
| 12.1. Definition                                                             | 30 |
| 12.2. Wie erhält man $f^{-1}$ aus $f$ ?                                      | 31 |
| 12.3. Ableitung der Umkehrfunktion                                           | 31 |
| 13. Die Arcusfunktionen                                                      | 31 |
| 13.1. Arcussinus                                                             | 31 |
| 13.2. Arcuscosinus                                                           | 31 |
| 13.3. Arcustangens                                                           | 32 |
| 13.4. Arcuscotangens                                                         | 32 |
| 14. Komplexe Zahlen                                                          | 32 |
| 14.1. Definition                                                             | 32 |
| 14.2. Addition und Multiplikation komplexer Zahlen                           | 32 |
| 14.3. Division komplexer Zahlen                                              | 33 |
| 14.4. Weitere Regeln für komplexe Zahlen                                     | 34 |
| 14.5. Fundamentalsatz der Algebra                                            | 34 |
| 14.6. Darstellungsformen komplexer Zahlen in der Gaußschen Zahlenebene       | 36 |
| 14.7. Multiplikation und Division komplexer Zahlen in trigonometrischer Form | 38 |
| 14.8. Betrag einer komplexen Zahl                                            | 39 |
| 14.9. Wurzeln aus komplexen Zahlen                                           | 40 |
| 14.10. komplexe Wurzeln aus 1                                                | 40 |
| 15. Matrizen                                                                 | 41 |
| 15.1. Definition einer Matrix                                                | 41 |
| 15.2. Verknüpfungen von Matrizen                                             | 41 |
| 15.3. Reguläre und inverse Matrizen                                          | 42 |
| 15.4. Transformation                                                         | 44 |
| 16. Determinanten                                                            | 46 |
| 16.1. Motivation                                                             | 46 |
| 16.2. Ermittlung von Determinanten                                           | 46 |
| 16.3. Rechenregeln für Determinanten                                         | 48 |
| 17. Vektoren                                                                 | 49 |
| 17.1. Vektoren                                                               | 49 |
| 17.2. Punkte                                                                 | 51 |
| 18. Skalarprodukt                                                            | 51 |
| 19. Vektorprodukt                                                            | 52 |
| 19.1. Die Gesetze des Vektorproduktes                                        | 53 |
| 20. Spatprodukt                                                              | 53 |
| 21. Geraden                                                                  | 53 |
| 21.1. Mittelpunkt einer Strecke                                              | 53 |
| 21.2. Parameterdarstellung                                                   | 53 |
| 21.3. Normalenform                                                           | 53 |

|                                                                        |    |
|------------------------------------------------------------------------|----|
| 21.4. Geraden–Konstellationen                                          | 53 |
| 21.5. Die Hesse-Gleichung                                              | 54 |
| 21.6. Abstand zwischen Punkt und Gerade                                | 54 |
| 21.7. Abstand zwischen zwei Geraden                                    | 55 |
| 21.8. Durchstoßpunkte mit den Grundebenen                              | 56 |
| 22. Ebenen                                                             | 56 |
| 22.1. Schwerpunkt eines Dreiecks                                       | 56 |
| 22.2. Ebenengleichung                                                  | 56 |
| 22.3. Abstand zwischen Punkt und Ebene                                 | 56 |
| 22.4. Abstand zwischen Gerade und Ebene                                | 57 |
| 22.5. Abstand zwischen zwei Ebenen                                     | 57 |
| 22.6. Schnitt zwischen Gerade und Ebene                                | 57 |
| 22.7. Schnitt zweier Ebenen                                            | 57 |
| 22.8. Spurpunkte                                                       | 57 |
| 22.9. Spurgeraden                                                      | 57 |
| 23. Spiegelungen                                                       | 57 |
| 23.1. Punkte                                                           | 57 |
| 23.2. Geraden                                                          | 58 |
| 23.3. Ebenen                                                           | 58 |
| 24. Kreise und Kugeln                                                  | 59 |
| 24.1. Kreis– und Kugelgleichung                                        | 59 |
| 24.2. Kreis                                                            | 59 |
| 24.3. Kugel                                                            | 59 |
| 24.4. Tangenten und Tangentialebenen                                   | 59 |
| 25. Folgen und Reihen                                                  | 59 |
| 26. Differentialrechnung                                               | 60 |
| 26.1. Symmetrie                                                        | 60 |
| 26.2. Verhalten                                                        | 61 |
| 26.3. Polstellen                                                       | 61 |
| 26.4. Asymptoten                                                       | 61 |
| 26.5. Nullstellen                                                      | 62 |
| 26.6. Ableitungen                                                      | 62 |
| 26.7. Tangenten                                                        | 63 |
| 26.8. Übersicht über bekannte Ableitungen                              | 63 |
| 26.9. Extremstellen                                                    | 63 |
| 26.10. Wendestellen                                                    | 64 |
| 26.11. Ganzrationale Funktionen                                        | 64 |
| 26.12. Komplette Funktionsuntersuchung ganzrationaler Funktionen       | 64 |
| 26.13. Gebrochenrationale Funktionen                                   | 65 |
| 26.14. Komplette Funktionsuntersuchung gebrochenrationaler Funktionen  | 65 |
| 26.15. Exponentialfunktionen                                           | 65 |
| 26.16. Logarithmusfunktionen                                           | 65 |
| 27. Integralrechnung                                                   | 66 |
| 27.1. Allgemeines                                                      | 66 |
| 27.2. Produktintegration                                               | 66 |
| 27.3. Tipps und Tricks zum Integrieren                                 | 66 |
| 27.4. Integration durch Substitution                                   | 66 |
| 27.5. Übersicht über bekannte Integrale und Stammfunktionen            | 67 |
| 27.6. Flächenberechnung                                                | 67 |
| 27.7. Volumenberechnung                                                | 68 |
| 27.8. Länge eines Bogens                                               | 68 |
| 28. Prof. Löffler: Abbildungen von Mengen (»Kommunikation von Mengen«) | 68 |
| 28.1. Komposition (Zusammensetzung) von Abbildungen                    | 69 |
| 29. Prof. Löffler: Zahlen                                              | 71 |
| 29.1. Zahlenräume                                                      | 71 |
| 29.2. Darstellung von Zahlen                                           | 71 |
| 29.3. Umrechnung zwischen beliebigen Zahlensystemen                    | 72 |
| 29.4. Reelle Zahlen: Zahlen zum Messen                                 | 73 |
| 29.5. Anwendung der Binärdarstellung                                   | 74 |
| 29.6. Effektive Berechnung der Quadratwurzel durch das Heron-Verfahren | 76 |

|                                                                        |     |
|------------------------------------------------------------------------|-----|
| 29.7. Zahldarstellung im Computer: Datentypen int, long, float, double | 76  |
| 29.8. Aufbau der Zahlen                                                | 77  |
| 30. Prof. Löffler: Polynome                                            | 81  |
| 30.1. Auswerten eines Polynoms an der Stelle $\beta$                   | 82  |
| 30.2. Nullstellen von Polynomen                                        | 84  |
| 30.3. Die allgemeine binomische Formel                                 | 84  |
| 31. Übungsaufgaben                                                     | 87  |
| 31.1. Einige Bezeichnungen                                             | 87  |
| 31.2. Grundbegriffe der Mengenlehre                                    | 87  |
| 31.3. Aussagenlogik                                                    | 87  |
| 31.4. Beweismethoden                                                   | 89  |
| 31.5. Abbildungen                                                      | 91  |
| 31.6. Permutation                                                      | 93  |
| 31.7. Zahlen                                                           | 93  |
| 31.8. Mächtigkeit von Zahlenmengen                                     | 94  |
| 31.9. Vollständige Induktion                                           | 94  |
| 31.10. Binominalkoeffizienten                                          | 97  |
| 31.11. Winkelfunktionen                                                | 97  |
| 31.12. Umkehrfunktion                                                  | 98  |
| 31.13. Die Arcusfunktionen                                             | 98  |
| 31.14. Komplexe Zahlen                                                 | 98  |
| 31.15. Matrizen                                                        | 100 |
| 31.16. Determinanten                                                   | 102 |
| 31.17. Vektoren                                                        | 102 |
| 31.18. Skalarprodukt                                                   | 102 |
| 31.19. Vektorprodukt                                                   | 102 |
| 31.20. Spatprodukt                                                     | 103 |
| 31.21. Geraden                                                         | 103 |
| 31.22. Ebenen                                                          | 104 |
| 31.23. Spiegelungen                                                    | 107 |
| 31.24. Kreise und Kugeln                                               | 107 |
| 31.25. Folgen und Reihen                                               | 107 |
| 31.26. Differentialrechnung                                            | 108 |
| 31.27. Integralrechnung                                                | 112 |
| Literatur                                                              | 112 |

## ABBILDUNGSVERZEICHNIS

|                                                                             |    |
|-----------------------------------------------------------------------------|----|
| 1 Zahlenstrahl                                                              | 7  |
| 2 $f(x) = x^2$ als Darstellung von $\{(x; x^2) \mid x \in \mathbb{R}\}$ .   | 7  |
| 3 Mengendiagramm                                                            | 8  |
| 4 Schnittmenge $A \cap B$                                                   | 9  |
| 5 Vereinigungsmenge $A \cup B$                                              | 10 |
| 6 Differenzmenge $A \setminus B$                                            | 10 |
| 7 Komplement $\overline{A}$                                                 | 10 |
| 8 Abbildungsvorschriften für surjektive, injektive und bijektive Abbildung. | 17 |
| 9 Erstes Diagonalverfahren                                                  | 19 |
| 10 Übersicht über die Teilmengen von $\mathbb{R}$                           | 20 |
| 11 Übliche Bezeichnungen am Dreieck                                         | 27 |
| 12 Winkelfunktionen am Einheitskreis                                        | 28 |
| 13 Achsensymmetrie der Sinusfunktion                                        | 28 |
| 14 Punktsymmetrie der Sinusfunktion                                         | 28 |
| 15 Der Einheitskreis                                                        | 29 |
| 16 Abtragung von Sinuswerten am Einheitskreis                               | 29 |
| 17 Konstruktion der Tangensfunktion am Einheitskreis                        | 30 |
| 18 Funktion und Umkehrfunktion                                              | 31 |

|                                                                            |    |
|----------------------------------------------------------------------------|----|
| 19 Die Funktionen $\sin x$ und $\arcsin x$                                 | 31 |
| 20 Die Funktionen $\cos x$ und $\arccos x$                                 | 32 |
| 21 Die Funktionen $\tan x$ und $\arctan x$                                 | 32 |
| 22 Die Funktionen $\cot x$ und $\operatorname{arccot} x$                   | 32 |
| 23 Darstellung von $z = 2 + 3i$ in der Gauß-Ebene und Addition             | 36 |
| 24 Subtraktion $z_1 - z_2 = z_1 + (-z_2)$                                  | 36 |
| 25 Polarform komplexer Zahlen                                              | 36 |
| 26 Zu Umwandlung von kartesischer in trigonometrische Form                 | 37 |
| 27 Zur trigonometrischen Form von $z = \frac{3}{2} - i\frac{1}{2}\sqrt{3}$ | 37 |
| 28 Differenz zweier Vektoren                                               | 50 |
| 29 Lotfußverfahren                                                         | 55 |
| 30 Spiegelung Punkt an Gerade in $\mathbb{R}^3$                            | 58 |
| 31 Venn-Diagramm zu $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$       | 88 |
| 32 Venn-Diagramm zu $\overline{A \cup B} = \overline{A} \cap \overline{B}$ | 88 |
| 33 Venn-Diagramm zu $\overline{A \cap B} = \overline{A} \cup \overline{B}$ | 88 |

## TABELLENVERZEICHNIS

|                            |    |
|----------------------------|----|
| 1 Bekannte Ableitungen     | 64 |
| 2 Bekannte Stammfunktionen | 67 |

### 1. EINIGE BEZEICHNUNGEN

| Zeichen            | Name                      | Anwendung; Sprechweise                                                                                                        |
|--------------------|---------------------------|-------------------------------------------------------------------------------------------------------------------------------|
| $\Rightarrow$      | Implikation               | $A \Rightarrow B$ ; aus der Aussage $A$ folgt die Aussage $B$ .                                                               |
| $\Leftrightarrow$  | Äquivalenz                | $A \Leftrightarrow B$ ; aus der Aussage $A$ folgt die Aussage $B$ und umgekehrt; $A$ gilt genau dann wenn $B$ gilt            |
| $:\Leftrightarrow$ | Äquivalenz per Definition | $A :\Leftrightarrow B$ ; es wird vereinbart, dass die Aussage $A$ anstelle der Aussage $B$ verwendet werden darf              |
| $:=$               | Gleichheit per Definition | $\alpha := \beta$ ; es wird vereinbart, dass die Bezeichnung $\alpha$ anstelle der Bezeichnung $\beta$ verwendet werden darf. |

Die »Äquivalenz per Definition« wird zur Kürzung zu langer Aussagen verwendet.

**Example 1.** Beispiele zu Bezeichnungen

**Example.** Seien  $x, y, a$  ganze Zahlen. Dann ist:

- $x + 2 = y \Rightarrow x < y$
- $a$  ist durch 6 teilbar  $\Leftrightarrow a$  ist durch 2 und 3 teilbar.
- $x$  ist gerade Zahl  $:\Leftrightarrow$  es gibt eine ganze Zahl  $k$  mit  $x = 2 \cdot k$ .
- $\operatorname{ld} x := \log_2 x$
- $\operatorname{lg} x := \log_{10} x$
- $\pi := \frac{\text{Kreisumfang}}{2 \cdot \text{Radius}}$

### 2. GRUNDBEGRIFFE DER MENGENLEHRE

Mengen sind die Sprache, in der in der Mathematik alles ausgedrückt wird. Was hier behandelt wird, ist nur die sog. »naive Mengenlehre«.

**Definition 1.** Menge nach Cantor (1845-1918)

**Definition.** Eine Menge  $M$  ist eine Zusammenfassung von bestimmten wohlunterschiedenen Objekten, den Elementen von  $M$ .

In einer Menge sind im Gegensatz zu Folgen die Elemente paarweise verschieden, d.h. mehrfach vorkommende Elemente werden gemäß Vereinbarung nur einfach gezählt. Will man mehrere gleiche Elemente in einer Menge unterbringen, muss man sie mit einer unterscheidbaren Zusatzbezeichnung versehen. Die Reihenfolge der Elemente in der Menge spielt keine Rolle:

$$\{1, 2, 3\} = \{3, 2, 1\} = \{1, 1, 2, 3\}$$

Mengen werden häufig mit Großbuchstaben bezeichnet:  $A, M, N, \Delta, \Omega$ . Die Symbole für Elemente sind frei wählbar, z.B.  $\{I, II, III\}$ .

**Example 2.** Beispiele für Mengen, spezielle Mengen

**Example.**

- Die Menge der Großbuchstaben.

$$U = UPPER = \{A, B, C, \dots, Z\}$$

- Die Menge der Kleinbuchstaben.

$$L = LOWER = \{a, b, c, \dots, z\}$$

- Die Menge der Ziffern.

$$D = DIGIT = \{0, 1, 2, \dots, 9\}$$

- Die Menge der Symbole für Hexadezimalzahlen.

$$X = XDIGIT = \{0, 1, 2, \dots, 9, A, B, \dots, F, a, b, \dots, f\}$$

- Die Menge der natürlichen Zahlen.

$$\mathbb{N} = \{1, 2, 3, 4, \dots\}$$

$$\mathbb{N}_0 = \{0, 1, 2, 3, 4, \dots\}$$

Die Menge der natürlichen Zahlen ist wichtig, weil die Elemente eine Anordnung haben; sie eignet sich zur Anordnung von anderen Elementen. Die natürlichen Zahlen reichen aus, um zu zählen, jedoch nicht, um zu messen.

- $\Delta_n = \{0, 1, 2, 3, \dots, n\}$ , z.B. die Menge  $\Delta_4 = \{0, 1, 2, 3, 4\}$ .
- Die Menge der ganzen Zahlen.

$$\mathbb{Z} = \{0, \pm 1, \pm 2, \dots\}$$

Es dauerte bis ins 18. Jahrhundert, bis die Existenz der ganzen Zahlen anerkannt wurde, d.h. länger als die Durchsetzung der komplexen Zahlen; denn: die Vorstellung von Zahlen »weniger als nichts« war ungewohnt.

- Die Menge der rationalen Zahlen.

$$\mathbb{Q} = \left\{ \frac{p}{q} \mid p, q \in \mathbb{Z}; q \neq 0 \right\}$$

- Die Menge der reellen Zahlen.

$$\mathbb{R}$$

- Menge der komplexen Zahlen

$$\mathbb{C}$$

- Eine Menge von Mengen.

$$\{\{1, 2\}, \{1, 3\}\}$$

Eine Menge mit zwei Elementen, von denen jedes wieder eine Menge mit zwei Elementen ist.

- Die leere Menge.

$$\emptyset := \{\}$$

Das Symbol  $\emptyset$  ist die international übliche Bezeichnung für die leere Menge. Die leere Menge gibt es nur einmal.

**2.1. Möglichkeiten, eine Menge anzugeben.** Dass ein Element in einer Menge enthalten ist, schreibt man als  $a \in M$ :  $a$  ist Element von  $M$ , die Negation ist  $a \notin M$ . Die Ausdrücke können auch umgedreht werden, und bleiben äquivalent:  $M \ni a$ .

**2.1.1. Aufzählen der Elemente.** Man unterscheidet zwischen endlichen (z.B.  $M = \{a, b, c\}$ ,  $\Delta_4 = \{0, 1, 2, 3, 4\}$ ) und unendlichen (z.B.  $\mathbb{N}$ ) Mengen. Unendliche Mengen können aufzählend mit Fortsetzungspunkten angegeben werden, z.B. ist  $P = \{1, 2, 3, 5, 7, 11, \dots\}$  die Menge aller Primzahlen. Diese Möglichkeit erschöpft sich, wenn die Folge der Elemente nicht offensichtlich erkennbar ist.

### 2.1.2. Angabe der Elementeigenschaften.

**Example 3.** Mengen definieren durch Angabe der gegenüber eine Obermenge einschränkenden Eigenschaften.

**Example.** Die Menge der natürlichen Zahlen von 1 bis 4:

$$N = \{n \in \mathbb{N} \mid n \leq 4\} = \{1, 2, 3, 4\}$$

Die Menge der Primzahlen:

$$P = \{p \in \mathbb{N} \mid p \text{ ist durch sich selbst und 1 teilbar}\}$$

Die Menge der Quadratzahlen: Der erste Teil des Ausdrucks ( $x^2$ ) gibt dabei immer die Gestalt der in der Menge enthaltenen Elemente an, der zweite Teil ( $x \in \mathbb{R}$ ) macht eine Aussage über Bedingungen, die für die Elemente gelten müssen.

$$Q = \{x^2 \mid x \in \mathbb{R}\}$$

$$P = \{x \in \mathbb{Z} \mid x > 0\}$$

$$\Delta_n = \{x \in \mathbb{N} \mid x \leq n\}$$

$$M = \{x \in \mathbb{Z} \mid x^2 = 4\} = \{-2, +2\}$$

$$K = \{x \in \mathbb{Z} \mid x^2 = 2\} = \{\}$$

$$M := \{x \mid x \text{ natürliche Zahl; } x \text{ ungerade; } 1 \leq x \leq 7\}$$

**2.2. Möglichkeiten, eine Menge darzustellen.** Dieses Unterkapitel wurde in der Vorlesung bei Prof. Löffler behandelt, bei Prof. Eichner jedoch nicht.

2.2.1. *Auf dem Zahlenstrahl.* Siehe Abbildung 1.



ABBILDUNG 1. Zahlenstrahl

2.2.2. *Graphen.* Ein Funktionsgraph gibt die Menge  $\{(x; f(x)) \mid x \in \mathbb{R}\}$  an. Zum Beispiel in 2 die Menge  $\{(x; x^2) \mid x \in \mathbb{R}\}$ .

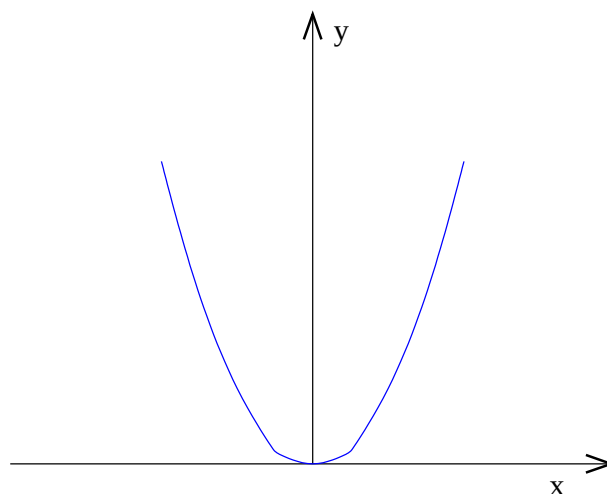


ABBILDUNG 2.  $f(x) = x^2$  als Darstellung von  $\{(x; x^2) \mid x \in \mathbb{R}\}$ .

2.2.3. *Mengendiagramme (Venn-Diagramme).* Sie dienen der Veranschaulichung von Mengenverknüpfungen. Siehe Abbildung 3.

**2.3. Relationen von Mengen (Beziehungen zwischen Mengen).**

**2.3.1. Gleichheit.** Enthalten zwei Mengen die gleichen Elemente, so sind sie gleich. Man schreibt  $M = N$ . Man prüft dies, indem man für jedes  $e \in M$  prüft, ob es auch  $e \in N$  ist (wenn ja: es gilt  $M \subseteq N$ ) und umgekehrt für jedes  $f \in N$  prüft, ob es auch  $f \in M$  ist (wenn ja: es gilt  $N \subseteq M$ ).

Für zwei Mengen  $\emptyset_1 = \{\}$ ,  $\emptyset_2 = \{\}$  gilt:  $\emptyset_1 \subset \emptyset_2$  und  $\emptyset_2 \subset \emptyset_1$ , also  $\emptyset_1 = \emptyset_2$ . Daraus folgt: alle leeren Mengen sind gleich; es gibt nur eine (»die«) leere Menge.

**Example 4.** Gleichheit von Mengen.

**Example.**

$$\{1; 2; 3\} = \{3; 2; 1\}$$

$$\{n \in \mathbb{N} \mid n \leq 4\} = \{1, 2, 3, 4\} = \{n \in \mathbb{N} \mid n < 5\}$$

$$\{x \in \mathbb{R} \mid x^2 - 1 = 0\} = \{-1; 1\}$$

**2.3.2. Obermenge und Inklusion / Teilmenge.**

**Definition 2.** Teilmenge, Obermenge, Gleichheit

**Definition.** Seien  $A, B$  Mengen. Dann gilt:

- $A \subseteq B : \Leftrightarrow$  Jedes Element von  $A$  ist auch Element von  $B$ . Man spricht: » $A$  ist Teilmenge von  $B$ «.
- $B \supseteq A : \Leftrightarrow A \subseteq B$ . Man spricht: » $B$  ist Obermenge von  $A$ «.
- $A = B : \Leftrightarrow (A \subseteq B \wedge B \subseteq A)$ . Man spricht: » $A$  gleich  $B$ «.

Merke:  $\emptyset \subset A$  für jede Menge  $A$  - denn man kann kein Element aus  $\emptyset$  finden, das nicht in  $A$  enthalten ist.

Wenn alle Elemente einer Menge  $A$  in einer Menge  $B$  enthalten sind (mathematisch  $e \in A \Rightarrow e \in B$ ), so ist  $A$  enthalten in  $B$ . Man schreibt:

- Entweder:  $A \subseteq B$ :  $A$  ist Teilmenge von oder gleich  $B$ <sup>1</sup>.
- Oder:  $A \subset B$ :  $A$  ist echte Teilmenge von  $B$ , also auch  $A \neq B$ <sup>2</sup>. Äquivalent ist die Schreibweise  $B \supset A$ :  $B$  enthält  $A$  bzw.  $B$  ist Obermenge von  $A$ .

Die Definition einer Teilmenge durch einschränkende Eigenschaften gegenüber einer Obermenge dient häufig zur Angabe von Mengen: Obermenge  $\Omega$ ,  $A = \{x \in \Omega \mid x \text{ erfüllt } \dots\}$ ; siehe 2.1.2.

**Example 5.** Beispiele für Teilmengen.

**Example.**

$$\{1\} \subseteq \{1, 2, 3\}$$

$$\{n \in \mathbb{N} \mid n \leq 7\} \subseteq \{n \in \mathbb{N} \mid n \leq 15\}$$

$$\mathbb{N} \subset \mathbb{Z}$$

**2.3.3. Potenzmenge.** Die Potenzmenge ist die Menge aller Teilmengen einer Menge  $A$ . Schreibweise:  $\wp(A)$ . Die Potenzmenge einer Menge mit  $x$  Elementen enthält  $2^x$  Elemente. Beispiel:  $A = \{1, 2\}$ ,  $\wp(A) = \{\emptyset, \{1\}, \{2\}, \{1, 2\}\}$ .

<sup>1</sup>Vergleiche das Zeichen  $\leq$ :  $1 \leq 7$  ist eine wahre Aussage.

<sup>2</sup>Vergleiche das Zeichen  $<$ .

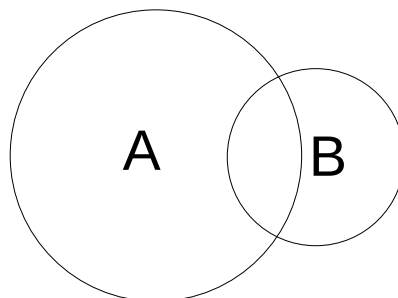


ABBILDUNG 3. Mengendiagramm



## 2.3.4. Mächtigkeit einer Menge.

**Definition 3.** Mächtigkeit einer Menge

**Definition.**  $|A| :\Leftrightarrow$  Anzahl der Elemente von  $A$ . Man liest: »Mächtigkeit der Menge  $A$ «. Diese Definition gilt vorläufig bis zur Betrachtung der Mächtigkeit unendlicher Mengen.

**Example 6.** Mächtigkeit einer Menge

**Example.**

$$|\emptyset| = 0$$

$$|\{1, 7, 5, 3\}| = 4$$

## 2.4. Verknüpfungen von Mengen.

## 2.4.1. Schnitt.

**Definition 4.** Schnittmenge

**Definition.** Sei  $M$  eine Menge und  $A, B \subset M$ . So wird die Schnittmenge (auch Durchschnitt) definiert als:  $A \cap B = \{x \in M \mid x \in A \wedge x \in B\}$ . Die Schnittmenge kann auch  $\emptyset$  sein.

**Example 7.** Schnittmenge.

**Example.** (siehe auch Abbildung 4)

$$A = \{1, 2, 3\}$$

$$B = \{1, 4, 5\}$$

$$A \cap B = \{1\}$$

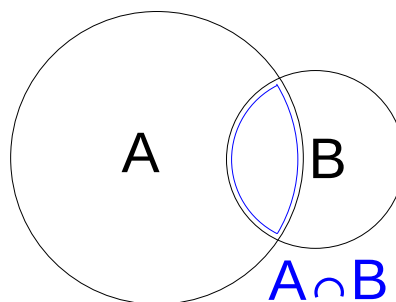


ABBILDUNG 4. Schnittmenge  $A \cap B$

## 2.4.2. Vereinigung.

**Definition 5.** Vereinigung

**Definition.** Sei  $M$  eine Menge und  $A, B \subset M$ . So wird die Vereinigungsmenge definiert als:  $A \cup B = \{x \in M \mid x \in A \vee x \in B\}$ <sup>3</sup>.

**Example 8.** Vereinigungsmenge.

**Example.** (siehe auch Abbildung 5)

$$A = \{1, 2, 3\}$$

$$B = \{1, 4, 5\}$$

$$A \cup B = \{1, 2, 3, 4, 5\}$$

<sup>3</sup>Das mathematische ODER ( $\vee$ ), wie es auch bei der Programmierung verwendet wird, ist kein exklusives ENTWEDER-ODER (wie manchmal in der Umgangssprache), sondern ein inklusives ODER AUCH.  $A \cup B$  enthält also auch Elemente, die sowohl in  $A$  als auch in  $B$  vorkommen.

### 2.4.3. Differenzmenge.

**Definition 6.** Differenzmenge » $A$  ohne  $B$ « (auch » $A$  minus  $B$ «)

**Definition.**  $A \setminus B := \{x \mid x \in A \wedge x \notin B\}$ . Diese Menge nennt man auch Restmenge.

**Example 9.** Differenzmenge.

**Example.** (siehe auch Abbildung 6)

$$\begin{aligned} A &= \{1, 2, 3\} \\ B &= \{1, 4, 5\} \\ A \setminus B &= \{2, 3\} \end{aligned}$$

2.4.4. *Komplement von  $A$ .* Ist  $A$  Teilmenge einer großen Grundmenge, so ist das Komplement der Menge  $A$  die Menge  $\bar{A} = \{x \mid x \notin A\}$ , manchmal auch geschrieben als  $\complement A$  oder  $A^C$ .

**Example 10.** Komplement.

**Example.** (siehe auch Abbildung 7)

$$\begin{aligned} A &= \{n \in \mathbb{N} \mid n \leq 4\} \\ \bar{A} &= \{n \in \mathbb{N} \mid n \geq 5\} \end{aligned}$$

2.5. **Supersatz.** Der Supersatz über Mengen ist ein Satz aus mehreren Teilen. Er beinhaltet Rechenregeln für Mengen.

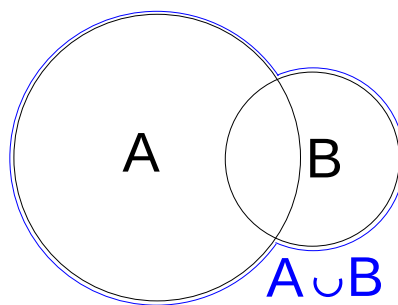


ABBILDUNG 5. Vereinigungsmenge  $A \cup B$

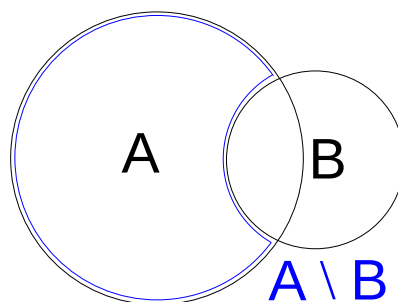


ABBILDUNG 6. Differenzmenge  $A \setminus B$

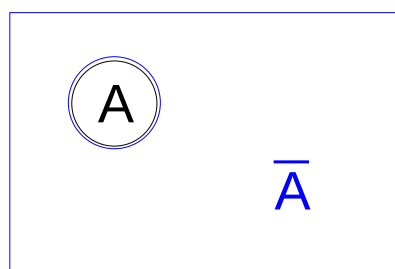


ABBILDUNG 7. Komplement  $\bar{A}$

## 2.5.1. Assoziativgesetze.

$$(1) \quad (A \cup B) \cup C = A \cup (B \cup C)$$

$$(2) \quad (A \cap B) \cap C = A \cap (B \cap C)$$

Die Reihenfolge bei der Verknüpfung von Mengen in Termen mit mehreren Verknüpfungen muss durch Klammern um je zwei Elemente definiert werden. Nach den Assoziativgesetzen (Gleichungen 1, 2) können Klammern entfallen, wenn nur Vereinigungsmengen oder nur Schnittmengen gebildet werden, denn die Reihenfolge der Operationen ist hier irrelevant. Von dieser Möglichkeit sollte man jedoch keinen Gebrauch machen.

## 2.5.2. Kommutativgesetze.

$$A \cap B = B \cap A$$

$$A \cup B = B \cup A$$

## 2.5.3. Distributivgesetze.

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

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

In Kombination mit dem Kommutativgesetz ergeben sich noch:

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

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

## 2.5.4. weitere Gesetze. Erstes Gesetz von DeMorgan:

$$\overline{A \cup B} = \overline{A} \cap \overline{B}$$

Zweites Gesetz von DeMorgan:

$$\overline{A \cap B} = \overline{A} \cup \overline{B}$$

Bei Anwendung der Gesetze von DeMorgan in Teilausdrücken ist das Ergebnis in eine Klammer zu setzen, denn eine Komplementbildung wie im Ausgangsterm impliziert stets eine Klammer.

Weitere Gesetze:

$$\overline{\overline{A}} = A$$

$$A \cap A = A$$

$$A \cup A = A$$

$$A \cap \emptyset = \emptyset$$

$$A \cup \emptyset = A$$

$$A \subset B \Rightarrow A \cap B = A$$

Priorität der Operationen:

(1) Klammerung

(2) Negierung. Jede Negierung impliziert eine Klammerung:  $\overline{A \cap B} = \overline{(A \cap B)}$

(3) Schnitt- und Vereinigungsmenge.

## 3. AUSSAGENLOGIK

## 3.1. Relationen.

**Definition 7.** aussagenlogische Variable**Definition.** Eine Variable, die nur die Werte  $f$  und  $w$  annehmen kann, heißt aussagenlogische Variable (AL-Variable). AL-Variable sind Platzhalter für beliebige Aussagen, die entweder falsch ( $f$ ) oder wahr ( $w$ ) sein können.**Definition 8.** Negation**Definition.** Sei eine Variable  $A$ .  $\overline{A}$  (»Nicht  $A$ «) ist definiert durch:

| $A$ | $\overline{A}$ |
|-----|----------------|
| $f$ | $w$            |
| $w$ | $f$            |

**Definition 9.** Und**Definition.** Seien zwei Variablen  $A, B$ .  $A \wedge B$  (» $A$  und  $B$ «) ist definiert durch:

| $A$ | $B$ | $A \wedge B$ |
|-----|-----|--------------|
| $f$ | $f$ | $f$          |
| $f$ | $w$ | $f$          |
| $w$ | $f$ | $f$          |
| $w$ | $w$ | $w$          |

**Definition 10.** Oder**Definition.** Seien zwei Variablen  $A, B$ .  $A \vee B$  (» $A$  oder  $B$ «) ist definiert durch:

| $A$ | $B$ | $A \vee B$ |
|-----|-----|------------|
| $f$ | $f$ | $f$        |
| $f$ | $w$ | $w$        |
| $w$ | $f$ | $w$        |
| $w$ | $w$ | $w$        |

**Definition 11.** Implikation**Definition.** Seien zwei Variablen  $A, B$ .  $A \rightarrow B$  (» $A$  impliziert  $B$ «; »wenn  $A$ , dann  $B$ «; »aus  $A$  folgt  $B$ «) ist definiert durch:

| $A$ | $B$ | $A \rightarrow B$ |
|-----|-----|-------------------|
| $f$ | $f$ | $w$               |
| $f$ | $w$ | $w$               |
| $w$ | $f$ | $f$               |
| $w$ | $w$ | $w$               |

**Definition 12.** Äquivalenz**Definition.** Seien zwei Variablen  $A, B$ .  $A \leftrightarrow B$  (» $A$  äquivalent  $B$ «; » $A$  genau dann, wenn  $B$ «) ist definiert durch:

| $A$ | $B$ | $A \leftrightarrow B$ |
|-----|-----|-----------------------|
| $f$ | $f$ | $w$                   |
| $f$ | $w$ | $f$                   |
| $w$ | $f$ | $f$                   |
| $w$ | $w$ | $w$                   |

Die Äquivalenz (Gleichwertigkeit) liefert  $w$ , wenn beide aussagenlogischen Variablen den gleichen Wahrheitswert haben.**Definition 13.** Exklusiv-Oder (Exor)**Definition.** Seien zwei Variablen  $A, B$ .  $A \oplus B$  (»entweder  $A$  oder  $B$ «) ist definiert durch:

| $A$ | $B$ | $A \oplus B$ |
|-----|-----|--------------|
| $f$ | $f$ | $f$          |
| $f$ | $w$ | $w$          |
| $w$ | $f$ | $w$          |
| $w$ | $w$ | $f$          |

Exor entspricht der eingeschränkten binären Addition (d.h. unter Vernachlässigung des Übertrags), deshalb das Zeichen  $\oplus$ . Exor ist die negierte Äquivalenz, deshalb auch das Zeichen  $\nleftrightarrow$ .

Alle Relationen der Aussagenlogik können durch UND, ODER, NICHT ausgedrückt werden, der Rest ist daher eigentlich unnötig.

**Definition 14.** rekursive Definition von Formeln der Aussagenlogik

**Definition.** Formeln der Aussagenlogik (AL) werden rekursiv definiert:

- (1) Jede AL-Variable  $A, B, C, \dots$  ist eine atomare Formel der Aussagenlogik.
- (2) Sind  $\alpha, \beta$  Formeln der Aussagenlogik, so auch  $\bar{\alpha}, (\alpha \wedge \beta), (\alpha \vee \beta), (\alpha \rightarrow \beta), (\alpha \leftrightarrow \beta), (\alpha \oplus \beta)$ .

**Example 11.** rekursive Definition von Formeln der Aussagenlogik

**Example.** Die Formel wird in mehreren Schritten aufgebaut:

- $A$
- $(\bar{A} \vee B)$
- $((\bar{A} \vee B) \leftrightarrow C)$

Vereinbarung über Klammern in Formeln der Aussagenlogik. Klammern können durch Vorrangregeln gespart werden. Bei gleicher Priorität müssen Klammern gesetzt werden. Die letzte, die gesamte Formel umfassende Klammer kann auch weggelassen werden. Prioritäten:

- (1) Negation
- (2) logisches Und  $\wedge$ , logisches Oder  $\vee$
- (3) Implikation  $\rightarrow$ , Äquivalenz  $\leftrightarrow$

**Example 12.** Weglassen von Klammern

**Example.**

$$\bar{A} \wedge B \rightarrow \overline{C \vee D}$$

steht für

$$\left( (\bar{A} \wedge B) \rightarrow \overline{(C \vee D)} \right)$$

**Definition 15.** Auswertung einer Formel  $\alpha$

**Definition.** Unter Auswertung einer Formel  $\alpha$  verstehen wir die Berechnung der Wahrheitswerte von  $\alpha$  für jede Belegung der Variablen in  $\alpha$  mit  $w$  und  $f$ .

**Example 13.** Auswertung der Formel  $\alpha := (A \vee B) \rightarrow (\bar{A} \leftrightarrow C)$

**Example.**

| $A$ | $B$ | $C$ | $(A \vee B)$ | $(\bar{A} \leftrightarrow C)$ | $\alpha := (A \vee B) \rightarrow (\bar{A} \leftrightarrow C)$ |
|-----|-----|-----|--------------|-------------------------------|----------------------------------------------------------------|
| $f$ | $f$ | $f$ | $f$          | $f$                           | $w$                                                            |
| $f$ | $f$ | $w$ | $f$          | $w$                           | $w$                                                            |
| $f$ | $w$ | $f$ | $w$          | $f$                           | $f$                                                            |
| $f$ | $w$ | $w$ | $w$          | $w$                           | $w$                                                            |
| $w$ | $f$ | $f$ | $w$          | $w$                           | $w$                                                            |
| $w$ | $f$ | $w$ | $w$          | $f$                           | $f$                                                            |
| $w$ | $w$ | $f$ | $w$          | $w$                           | $w$                                                            |
| $w$ | $w$ | $w$ | $w$          | $f$                           | $f$                                                            |

Anwendbar ist eine solche Auswertung einer Formel z.B. zur Definition komplizierter Bedingungen beim Programmieren, z.B. der Entscheidung, ob ein Tag ein Schalttag ist.

**Definition 16.** Werteverlauf einer Formel

**Definition.** Die Ergebnisse der Auswertung einer Formel sind ihr Werteverlauf.

**Definition 17.** Gleichwertigkeit von Formeln

**Definition.** Zwei Formeln  $\alpha, \beta$  heißen gleichwertig, in Zeichen  $\alpha = \beta$ , wenn ihre Auswertungen den gleichen Werteverlauf liefern.

**Example 14.** Gleichwertigkeit von Formeln: Implikation ausgedrückt durch Und und Oder

**Example.** Der Werteverlauf der Implikation ist:

| $A$ | $B$ | $A \rightarrow B$ |
|-----|-----|-------------------|
| $f$ | $f$ | $w$               |
| $f$ | $w$ | $w$               |
| $w$ | $f$ | $f$               |
| $w$ | $w$ | $w$               |

Ein anderer Werteverlauf ist diesem gleich:

| $A$ | $B$ | $A \vee B$ |
|-----|-----|------------|
| $f$ | $f$ | $w$        |
| $f$ | $w$ | $w$        |
| $w$ | $f$ | $f$        |
| $w$ | $w$ | $w$        |

Damit gilt folgende Gleichwertigkeit, die beim Programmieren, wo ja nur Und und Oder zur Verfügung stehen, ausgenutzt werden kann:

$$A \rightarrow B = \overline{A} \vee B$$

**Definition 18.** allgemeingültige Formel, Tautologie

**Definition.** Sei  $\phi$  eine Formel der Aussagenlogik.  $\phi$  heißt »allgemeingültig« oder eine »Tautologie«, wenn sie immer den Wert  $w$  liefert.

**Definition 19.** Kontradiktion

**Definition.** Sei  $\phi$  eine Formel der Aussagenlogik.  $\phi$  heißt »Kontradiktion«, wenn sie immer  $f$  liefert.

**Definition 20.** Gleichwertigkeit von Formeln der Aussagenlogik über Tautologie erkennen

**Definition.** Seien  $\alpha, \beta$  Formeln. Es gilt:  $\alpha = \beta \Leftrightarrow ((\alpha \leftrightarrow \beta)$  ist eine Tautologie)

**Example 15.** Gleichwertigkeit von Formeln der Aussagenlogik über Tautologie erkennen

**Example.**

$$\phi := (A \rightarrow B) \leftrightarrow (\overline{A} \vee B)$$

An folgender Wertetabelle kann man erkennen, dass  $\phi$  eine Tautologie ist:

| $(A \rightarrow B)$ | $(\overline{A} \vee B)$ | $\phi := (A \rightarrow B) \leftrightarrow (\overline{A} \vee B)$ |
|---------------------|-------------------------|-------------------------------------------------------------------|
| $w$                 | $w$                     | $w$                                                               |
| $w$                 | $w$                     | $w$                                                               |
| $f$                 | $f$                     | $w$                                                               |
| $w$                 | $w$                     | $w$                                                               |

Nach Definition 20 gilt nun, dass die beiden Teilformeln von  $\phi$  gleichwertig sind, wir also schreiben dürfen  $(A \rightarrow B) = (\overline{A} \vee B)$ .

Dieses Problem, Formeln schnell auf eine Tautologie zu prüfen, ist eines der großen ungelösten Probleme der theoretischen Informatik. Es ist ähnlich wie die Entschlüsselung immer prinzipiell, aber aus Zeitgründen nicht zu lösen.

**3.2. Grundgesetze der Aussagenlogik.** Dies ist eigentlich Stoff der Vorlesung »Mathematik 2«. Hier daher nur eine kurze und unvollständige Aufführung, beachte die Analogie zur Mengenlehre.

**3.2.1. Kommutativgesetze.**

$$\begin{aligned}\alpha \vee \beta &= \beta \vee \alpha \\ \alpha \wedge \beta &= \beta \wedge \alpha\end{aligned}$$

**3.2.2. Assoziativgesetze.**

$$\begin{aligned}\alpha \vee (\beta \vee \gamma) &= (\alpha \vee \beta) \vee \gamma \\ \alpha \wedge (\beta \wedge \gamma) &= (\alpha \wedge \beta) \wedge \gamma\end{aligned}$$

**3.2.3. DeMorgan.**

$$(3) \quad \begin{aligned}\overline{\alpha \vee \beta} &= \overline{\alpha} \wedge \overline{\beta} \\ \overline{\alpha \wedge \beta} &= \overline{\alpha} \vee \overline{\beta}\end{aligned}$$

**3.2.4. Idempotenz.**

$$\begin{aligned}\alpha \vee \alpha &= \alpha \\ \alpha \wedge \alpha &= \alpha\end{aligned}$$

#### 4. BEWEISMETHODEN

Es gibt nur drei Beweisverfahren in der Mathematik: direkter Beweis, indirekter Beweis, vollständige Induktion (siehe Kapitel 9) und ihre Verallgemeinerung. Wir betrachten mathematische Aussagen der Form  $A \Rightarrow B$  (»Aus der Aussage  $A$  folgt die Aussage  $B$ «).  $A$  heißt Voraussetzung,  $B$  heißt Behauptung der Aussage. Wir nehmen stets an, dass  $A$  wahr ist.

**4.1. Direkter Beweis einer Behauptung  $B$ .** Das bedeutet: Aus  $A$  wird durch eine Kette richtiger Folgerungen die Aussage  $B$  hergeleitet. Eine Rechnung in mehreren Schritten ist ein solcher direkter Beweis. Schema:

$$A \Rightarrow A_1 \Rightarrow A_2 \Rightarrow \dots \Rightarrow A_n \Rightarrow B$$

**4.2. Indirekter Beweis einer Behauptung  $B$ .** Das bedeutet: Man nimmt an,  $B$  sei falsch. Durch eine Kette richtiger Folgerungen leitet man einen Widerspruch zu  $A$  her. Da  $A$  wahr ist, muss die Annahme,  $B$  sei falsch, selbst falsch sein. Also ist  $B$  richtig. Schema:

$$\overline{B} \Rightarrow B_1 \Rightarrow B_2 \Rightarrow \dots \Rightarrow B_n \Rightarrow \overline{A}$$

das heißt  $\overline{A}$  steht im Widerspruch zu der richtigen Voraussetzung  $A$ , also gilt  $B$ .

### 4.3. Beispiele zu Beweismethoden.

**Definition 21.** Rest bei ganzzahliger Division

**Definition.** Seien zwei ganze Zahlen  $n \geq 0$  und  $m > 1$ . Dann gibt es genau ein Paar ganzer Zahlen  $k$  und  $r$  mit

$$n = mk + r \text{ mit } k \geq 0, 0 \leq r \leq m - 1$$

Hierbei heißt  $r$  der Rest der Division von  $n$  durch  $m$ . Man sagt,  $n$  ist durch  $m$  teilbar, wenn  $r = 0$ . Man schreibt auch  $r = n \bmod m$ ,  $0 \leq r \leq m - 1$ .

**Definition 22.** gerade Zahl, ungerade Zahl

**Definition.** Sei  $n \in \mathbb{Z}$ .

- (1)  $n$  heißt gerade  $:\Leftrightarrow$  es existiert ein  $k \in \mathbb{Z}$  mit  $n = 2k$ .
- (2)  $n$  heißt ungerade  $:\Leftrightarrow$  es existiert ein  $k \in \mathbb{Z}$  mit  $n = 2k + 1$  (das heißt,  $n$  ist nicht gerade).

Bemerkung:  $n$  ist also ungerade  $\Leftrightarrow n$  nicht gerade.

**Example 16.** Beweis von  $n$  gerade  $\Rightarrow n^2$  gerade

**Example.** Gegeben sei die Aussage  $n$  gerade  $\Rightarrow n^2$  gerade. Man kann nicht sofort erkennen, ob man etwas mit direktem oder indirektem Beweis beweisen muss, muss also probieren. Man hält sich bei direktem und indirektem Beweis ganz eng an die Definitionen: man ersetzt die Voraussetzung durch ihre Definition und fasst diese Terme oder Aussagen dann durch Operationen und Rückwärtsanwendung der Definitionen zur Behauptung zusammen.

$$\begin{aligned} & n \text{ gerade} \\ \Rightarrow & n = 2k \text{ (mit } k \in \mathbb{Z}) \\ \Rightarrow & n^2 = 4k^2 = 2(2k^2) \\ \Rightarrow & n^2 = 2k' \text{ (mit } k' \in \mathbb{Z}) \\ \Rightarrow & n^2 \text{ gerade} \end{aligned}$$

**Example 17.** Beweis von  $n^2$  gerade  $\Rightarrow n$  gerade

**Example.** Hier kommt man nur mit indirektem Beweis zum Ziel. Annahme, dass die Behauptung falsch ist:  $n$  nicht gerade, d.h.  $n$  ungerade. Aus dieser Annahme wollen wir nun durch direktes Schließen zu einem Widerspruch zur Voraussetzung » $n^2$  gerade« kommen:

$$\begin{aligned} & n \text{ ungerade} \\ \Rightarrow & n = 2k + 1 \\ \Rightarrow & n^2 = 4k^2 + 4k + 1 \\ \Rightarrow & n^2 = 2(2k^2 + 2k) + 1 \\ \Rightarrow & n^2 = 2k' + 1 \\ \Rightarrow & n^2 \text{ ungerade} \\ \Rightarrow & \text{Widerspruch zur Voraussetzung} \\ \Rightarrow & \text{Annahme falsch} \\ \Rightarrow & n \text{ gerade} \end{aligned}$$

**Definition 23.** Zusammenfassung der Beziehung von  $n$  und  $n^2$  als gerade Zahlen

**Definition.** Aus den in Beispielen 16 und 17 bewiesenen Behauptungen folgt:

$$n \text{ gerade} \Leftrightarrow n^2 \text{ gerade}$$

**Example 18.** Mehrteiliger Beweis

**Example.**

**Satz:** Seien  $a, b \in \mathbb{Z}$ . Es gilt

$$ab \text{ gerade} \Leftrightarrow a \text{ gerade} \vee b \text{ gerade}$$

**Beweis:** Der Satz enthält zwei einzelne Sätze, die getrennt bewiesen werden müssen:

- $ab \text{ gerade} \Rightarrow a \text{ gerade} \vee b \text{ gerade}$

Beweis mit indirektem Beweis. Die Annahme ist nun: » $a \text{ gerade} \vee b \text{ gerade}$ « ist falsch«, d.h. nach Gleichung 3: » $a \text{ ungerade} \wedge b \text{ ungerade}$ «.

$$\begin{aligned} & a \text{ ungerade} \wedge b \text{ ungerade} \\ \Rightarrow & a = 2k + 1 \wedge b = 2k' + 1 \\ \Rightarrow & ab = (2k + 1)(2k' + 1) \\ \Rightarrow & ab = 4kk' + 2k + 2k' + 1 \\ \Rightarrow & ab = 2(2kk' + k + k') + 1 \\ \Rightarrow & ab = 2k'' + 1 \\ \Rightarrow & ab \text{ ungerade} \\ \Rightarrow & \text{Widerspruch zur Voraussetzung} \\ \Rightarrow & \text{Annahme falsch} \\ \Rightarrow & a \text{ gerade} \vee b \text{ gerade} \end{aligned}$$

- $a \text{ gerade} \vee b \text{ gerade} \Rightarrow ab \text{ gerade}$

Beweis mit direktem Beweis. Die Voraussetzung hat durch das »Oder« 3 Fälle - für alle muss überprüft werden, ob die Behauptung richtig ist.

(1) Fall:

$$\begin{aligned} & a \text{ gerade} \wedge b \text{ ungerade} \\ \Rightarrow & a = 2k, b = 2k' + 1 \\ \Rightarrow & ab = 2k(2k' + 1) \\ \Rightarrow & ab = 2(2kk' + k) \\ \Rightarrow & ab = 2k'' \\ \Rightarrow & ab \text{ gerade} \end{aligned}$$

Das Zeichen » $\wedge$ « kann auch als »Und« oder wie im Folgenden als »,« geschrieben werden.

(2) Fall:

$$a \text{ ungerade}, b \text{ gerade}$$

Analog zum 1. Fall. In dieser Art die analoge Beweisführung dem Leser zu überlassen ist in Beweisen zulässig.

(3) Fall:

$$\begin{aligned} & a \text{ gerade}, b \text{ gerade} \\ \Rightarrow & a = 2k, b = 2k' \\ \Rightarrow & ab = 2k \cdot 2k' \\ \Rightarrow & ab = 2(2kk') \\ \Rightarrow & ab = 2k'' \\ \Rightarrow & ab \text{ gerade} \end{aligned}$$

## 5. ABBILDUNGEN

**Definition 24.** Abbildung, Definitionsbereich, Bildbereich

**Definition.** Seien  $A, B$  Mengen. Eine Abbildung  $f$  von  $A$  in  $B$  (in Zeichen:  $f : A \mapsto B$ ) ist eine Vorschrift, die jedem  $x \in A$  genau ein Element  $y \in B$  zuordnet (in Zeichen:  $y = f(x)$  oder  $f : x \mapsto y$ ).  $A$  heißt Definitionsbereich (oder: Quellbereich, Grundmenge) von  $f$ ,  $B$  heißt Bildbereich (oder: Zielbereich, Bildmenge) von  $f$ .

**Definition 25.** Abbildungen: weitere Bezeichnungen

**Definition.** Sei  $y = f(x)$ .  $x$  heißt unabhängige Variable (oder: Argument, Urbild von  $y$ ),  $y$  heißt abhängige Variable (oder: Funktionswert von  $x$  unter  $f$ , Bild von  $x$ ).

**Definition 26.** Bild von  $C$  unter  $f$ ,  $f$  eingeschränkt auf  $C$



**Definition.** Sei  $C \subset A$ ;  $f : A \mapsto B$ .

- (1)  $f(C) := \{f(x) \mid x \in C\}$  heißt Bild von  $C$  unter  $f$ .
- (2)  $f|_C : C \mapsto B$  mit  $f|_C(x) := f(x)$  für  $x \in C$  heißt Einschränkung (oder: Restriktion) von  $f$  auf  $C$ . Anstelle  $f|_C$  schreibt man auch  $f \mid C$ . Man liest: » $f$  eingeschränkt auf  $C$ «.

**Definition 27.** surjektive, injektive, bijektive Abbildung

**Definition.**  $f : A \mapsto B$  heißt (vgl. auch Abbildung 8):

- »surjektiv« oder »Abbildung von  $A$  auf  $B$ « :  $\Leftrightarrow B = f(A)$ . D.h.: Jedes  $y \in B$  ist Bild eines oder mehrerer  $x \in A$ ! Eine nicht surjektive Abbildung heißt »Abbildung von  $A$  in  $B$ «.
- »injektiv« oder »eineindeutige Abbildung von  $A$  in  $B$ « :  $\Leftrightarrow$  Für alle  $x_1, x_2 \in A$  gilt:  $(x_1 \neq x_2 \Leftrightarrow f(x_1) \neq f(x_2))$ . D.h.: Wenn die Urbilder verschieden sind, müssen auch die Bilder verschieden sein; also bei jedem Element von  $B$  kommt kein oder höchstens ein Pfeil an.
- »bijektiv« oder »umkehrbar eindeutig« :  $\Leftrightarrow f$  ist injektiv und surjektiv. D.h.:  $|A| = |B|$  - wenn es eine bijektive Abbildung  $f : A \mapsto B$  gibt, dann haben  $A$  und  $B$  gleich viele Elemente. Dies gilt auch für unendliche Mengen  $A, B$ .

**Definition 28.** Gleiche Mächtigkeit von Mengen über Bijektion

**Definition.** Zwei Mengen  $A, B$  heißen gleich mächtig, in Zeichen  $|A| = |B|$ , wenn es eine Bijektion  $f : A \mapsto B$  gibt.

**Definition 29.** inverse Abbildung

**Definition.** Sei  $f : A \mapsto B$  bijektiv. Dann heißt  $f^{-1} : B \mapsto A$  die inverse Abbildung (oder: Umkehrfunktion) von  $f$  und ist definiert durch  $f^{-1}(y) := x : \Leftrightarrow y = f(x)$ .

## 6. PERMUTATION

**Definition 30.** Permutation

**Definition.** Sei  $A$  eine endliche Menge. Sei  $n = |A|$ . Eine Permutation von  $A$  ist eine Bijektion  $p : \{1, 2, \dots, n\} \mapsto A$ .

**Example 19.** Permutation

**Example.**  $A = \{a, b, c\}$ . Dann sind sämtliche Permutationen von  $A$  in Form einer Tabelle:

|       | 1   | 2   | 3   |
|-------|-----|-----|-----|
| $p_1$ | $a$ | $b$ | $c$ |
| $p_2$ | $a$ | $c$ | $b$ |
| $p_3$ | $b$ | $a$ | $c$ |
| $p_4$ | $b$ | $c$ | $a$ |
| $p_5$ | $c$ | $a$ | $b$ |
| $p_6$ | $c$ | $b$ | $a$ |

Die Permutationen entsprechen sämtlichen Anordnungen der Elemente.

ABBILDUNG 8. Abbildungsvorschriften für surjektive, injektive und bijektive Abbildung.

## 7. ZAHLEN

### 7.1. Grundgesetze der reellen Zahlen.

**V:** Verknüpfungsaxiom

**A:** Assoziativgesetz  $(a + b) + c = a + (b + c)$ ,  $(a \cdot b) \cdot c = a \cdot (b \cdot c)$

**N:** Existenz eines neutralen Elementes  $0 + a = a + 0 = a$ ,  $1 \cdot a = a \cdot 1 = a$

**I:** Existenz des inversen Elementes  $a + (-a) = 0$ ,  $a \cdot a^{-1} = 1$  ( $a \neq 0$ )

**K:** Kommutativgesetz  $a + b = b + a$ ,  $a \cdot b = b \cdot a$

**7.2. Warum gilt  $\mathbb{Q} \subset \mathbb{R}$ ?** Zahlen aus  $\mathbb{Q}$  haben die Gestalt  $\frac{z}{n}$  mit  $z \in \mathbb{Z}, n \in \mathbb{N}$ ; Zahlen aus  $\mathbb{R}$  haben die Gestalt von Dezimalbrüchen. Wie wird  $\mathbb{Q}$  in  $\mathbb{R}$  eingebettet, d.h. wie kann man Brüche durch Dezimalbrüche ersetzen? Indem man die rationalen Zahlen in endliche oder periodische Dezimalbrüche umwandelt und umgekehrt. Nur die Elemente aus  $\mathbb{R}$ , die endlich oder periodisch sind, sind auch Elemente von  $\mathbb{Q}$ .

**Example 20.** Umwandlung von rationaler Zahl in eine reelle Zahl

**Example.** Durch schriftliche Division erhält man endliche oder periodische Dezimalbrüche:

- $\frac{5}{6} = 5 : 6 = 0,833\dots = 0,8\bar{3}$
- $\frac{3}{2} = 3 : 2 = 1,50\dots = 1,5$

**Example 21.** Umwandlung von reeller Zahl in rationale Zahl

**Example.** Man multipliziert die reelle Zahl einmal so, dass die erste Periode genau vor das Komma kommt, und einmal so, dass die erste Periode genau hinter das Komma kommt. Dann subtrahiert man beide Ergebnisse und erhält einen Bruch, der ggf. noch gekürzt werden kann.

•

$$\begin{aligned} x &= 0,23\bar{56} \\ 10000x &= 2356,\bar{56} \\ 100x &= 23,\bar{56} \\ \Rightarrow 9900x &= 2333 \\ &\quad \underline{2333} \\ \Leftrightarrow x &= \frac{9999}{9900} \end{aligned}$$

•

$$\begin{aligned} x &= 0,\bar{9} \\ 10x &= 9,\bar{9} \\ x &= 0,\bar{9} \\ \Rightarrow 9x &= 9 \\ &\quad \underline{9} \\ \Leftrightarrow x &= \frac{9}{9} = 1 \end{aligned}$$

Die Zahldarstellung der reellen Zahlen also ist nicht eindeutig! Allgemein kann jeder endliche Dezimalbruch durch einen periodischen Dezimalbruch dargestellt werden:  $5,47 = 5,46\bar{9}$ . Entfernt man so alle endlichen Dezimalbrüche, gibt es in  $\mathbb{R}$  keine zwei Zeichenketten mehr, die dieselbe Zahl darstellen und die Darstellung in  $\mathbb{R}$  ist nun eindeutig!

**7.3. Irrationale Zahlen.** Eine reelle Zahl gehört genau dann zu  $\mathbb{Q}$ , wenn sie als endlicher oder periodischer Dezimalbruch darstellbar ist. Es gibt also Zahlen, die nicht als Brüche darstellbar sind, d.h.  $\mathbb{R}$  enthält tatsächlich mehr Elemente als  $\mathbb{Q}$ .

**Example 22.** Eine Zahl nicht in  $\mathbb{Q}$  (indirekter Beweis)

**Example.** Die Gleichung  $x^2 = 2$  besitzt keine Lösung in  $\mathbb{Q}$ , d.h.  $\sqrt{2} \notin \mathbb{Q}$  (d.h.  $\sqrt{2}$  ist nicht als Bruch darstellbar). Beweis mit einer Variante des indirekten Beweises. Also Annahme:  $x$  sei ein Bruch, d.h.  $x = \frac{z}{n}$ . Wir benötigen noch eine Voraussetzung, zu der wir mit dieser Annahme einen Widerspruch erzeugen wollen; dazu verwendet man gerne Selbstverständlichkeiten, hier:  $\frac{z}{n}$  sei vollständig gekürzt, d.h.  $z$  und  $n$  besitzen nur die gemeinsamen Teiler 1 oder  $-1$ . Dass sich diese Voraussetzung auf unsere Annahme bezieht, ist die hier verwandte, zulässige Variation des indirekten Beweises.

Aus der ursprünglichen Annahme  $x = \frac{z}{n}$  folgt  $x^2 = \frac{z^2}{n^2}$ :

$$\begin{aligned} (4) \quad & \frac{z^2}{n^2} = 2 \\ (5) \quad & \Rightarrow z^2 = 2n^2 \\ (6) \quad & \Rightarrow z^2 \text{ ist gerade} \\ (7) \quad & \Rightarrow z \text{ ist gerade} \\ (8) \quad & \Rightarrow z = 2k \end{aligned}$$

Die Beziehung  $z^2$  ist gerade  $\Rightarrow z$  ist gerade wurde bereits früher bewiesen. Es gilt andererseits:

$$\begin{aligned} & \frac{z^2}{n^2} = 2 \\ \Rightarrow & n^2 = \frac{z^2}{2} \end{aligned}$$

Mit Gleichung 8 folgt:

$$\begin{aligned}
 &\Rightarrow n^2 = \frac{4k^2}{2} = 2k^2 \\
 &\Rightarrow n^2 \text{ gerade} \\
 &\Rightarrow n \text{ gerade} \\
 &\Rightarrow n = 2k' \\
 &\Rightarrow \frac{z}{n} \text{ mit 2 kürzbar} \\
 &\Rightarrow \text{Widerspruch zur Voraussetzung} \\
 &\Rightarrow x \text{ ist nicht als Bruch darstellbar}
 \end{aligned}$$

Also ist  $\sqrt{2}$  nicht als Bruch darstellbar. Also ist  $\mathbb{Q}$  eine echte Teilmenge von  $\mathbb{R}$ .

**Definition 31.** Echte Teilmenge

**Definition.**  $A$  ist eine echte Teilmenge von  $B$ , wenn  $A \subseteq B$ ,  $B \setminus A \neq \emptyset$ .

**Definition 32.** Irrationale Zahlen

**Definition.** Zahlen, die nicht rational sind, also nicht als Bruch darstellbar sind, heißen irrational.

#### 7.4. Algebraische und transzendente Zahlen.

**Definition 33.** algebraische Zahl

**Definition.** Eine reelle Zahl heißt algebraisch, wenn sie Lösung einer Gleichung der Form  $a_n x^n + a_{n-1} x^{n-1} + \dots + a_0 = 0$  mit  $a_i \in \mathbb{Q}$  (d.h. Nullstelle eines Polynoms) ist. Auch  $a_i \in \mathbb{Z}$  ist zulässig, da man die Zahlen aus  $\mathbb{Q}$  ja auf einen Hauptnenner bringen kann und die Gleichung mit diesem Hauptnenner multiplizieren kann. Die Lösungen dieser Polynome heißen Wurzeln.

**Example 23.** algebraische Zahl

**Example.**  $\sqrt{2}$  ist algebraisch, denn sie löst:  $x^2 - 2 = 0$ . Alle  $\frac{z}{n} \in \mathbb{Q}$  sind algebraische Zahlen, denn sie lösen  $x - \frac{z}{n} = 0$

**Definition 34.** transzendente Zahl

**Definition.** Eine reelle Zahl heißt transzendent, wenn sie nicht algebraisch ist. Zum Beispiel sind  $\pi$  und  $e$  transzendente Zahlen. Es gibt mehr transzendente als algebraische Zahlen (siehe später).

### 8. MÄCHTIGKEITEN VON ZAHLENMENGEN

**Definition 35.** Mächtigkeit einer Zahlenmenge (Wiederholung)

**Definition.** Zwei Mengen  $A, B$  heißen gleichmächtig, in Zeichen  $|A| = |B|$ , wenn es eine Bijektion von  $A$  auf  $B$  gibt.

**Definition 36.** Mächtigkeit von  $\mathbb{N}$  und  $\mathbb{Z}$

**Definition.** Es gilt  $|\mathbb{N}| = |\mathbb{Z}|$ . Man zählt alle Elemente von  $\mathbb{Z}$  in der Form:

$$\mathbb{Z} = \{0, 1, -1, 2, -2, 3, -3, \dots\}$$

Durch Nummerieren der Elemente erhält man nun die Bijektion:

$$f(z) = \begin{cases} 2z & \text{für } z \geq 0 \\ 2|z| + 1 & \text{für } z < 0 \end{cases}$$

Diese Funktion  $f: \mathbb{Z} \rightarrow \mathbb{N}$  ist eine Bijektion, d.h. sie ordnet jeder ganzen Zahl eine natürliche Zahl zu. Damit haben  $\mathbb{N}$  und  $\mathbb{Z}$  gleich viele Elemente (man sagt: sie sind gleich mächtig), obwohl  $\mathbb{N} \subset \mathbb{Z}$ .

**Definition 37.** Mächtigkeit von  $\mathbb{Q}$

**Definition.** Es gilt  $|\mathbb{N}| = |\mathbb{Q}|$ . Beweis durch das Erste Diagonalverfahren als Abzählverfahren (siehe Abbildung 9) - hierdurch hat man die Bijektion gefunden.

#### ABBILDUNG 9. Erstes Diagonalverfahren

**Definition 38.** Abzählbarkeit einer Menge

**Definition.** Eine Menge  $A$  heißt abzählbar, wenn sie endlich ist oder wenn  $|A| = |\mathbb{N}|$ . Damit sind  $\mathbb{N}$ ,  $\mathbb{Z}$ ,  $\mathbb{Q}$  abzählbar.

**Definition 39.** Überabzählbarkeit einer Menge

**Definition.** Eine Menge, die nicht abzählbar ist, heißt überabzählbar.

**Example 24.** Überabzählbarkeit von  $\mathbb{R}$

**Example.** Wir beweisen mit indirektem Beweis, dass die Zahlen  $x$  mit  $0 < x \leq 1$  nicht abzählbar sind. Jede dieser Zahlen ist auf genau eine Weise als unendlich langer Dezimalbruch darstellbar, wobei unendlich viele Ziffern  $\neq 0$  sind:

$$x = 0, x_1 x_2 x_3 \dots$$

Dazu werden ersetzt:  $1 = 0, \bar{9}$ ;  $0,476 = 0,475\bar{9}$ . Annahme zum indirekten Beweis: alle Zahlen  $x = 0, x_1 x_2 x_3 \dots$  mit  $0 < x \leq 1$  sind abzählbar. Also denken wir uns alle  $x$  nach Größe geordnet hingeschrieben:

$$\begin{aligned} a_1 &= 0, a_{11} a_{12} a_{13} a_{14} \dots \\ a_2 &= 0, a_{21} a_{22} a_{23} a_{24} \dots \\ a_3 &= 0, a_{31} a_{32} a_{33} a_{34} \dots \\ &\vdots \end{aligned}$$

Jetzt wird gezeigt, dass damit eine Zahl nicht getroffen wird: Man bilde  $b = 0, b_1 b_2 b_3 \dots$  wie folgt: Wähle  $b_1 \neq a_{11}$ ,  $b_2 \neq a_{22}$ ,  $b_3 \neq a_{33}$ ,  $\dots$ ; alle  $b_i$  können  $\neq 0$  gewählt werden. Es gilt damit:  $0 < b \leq 1$ , d.h.  $b$  muss unter den  $a_i$  vorkommen (da Zählbarkeit angenommen wurde). Dies ist unmöglich, denn wäre  $b = a_k$ , dann müsste  $b_k = a_{kk}$  sein, entgegen der Konstruktion von  $b$ ! Widerspruch zur Voraussetzung der Abzählbarkeit der Zahlen zwischen 0 und 1, also ist auch  $\mathbb{R}$  nicht abzählbar. Man schreibt  $|\mathbb{N}| < |\mathbb{R}|$ .

Gibt es Mengen mit noch mehr Elementen als in  $\mathbb{R}$ ? Ja: die Potenzmenge von  $\mathbb{R}$  hat noch mehr Elemente, höhere Potenzmengen jeweils noch mehr. Es ist unerheblich für die Mathematik, ob man eine Unendlichkeit zwischen  $\mathbb{N}$  und  $\mathbb{R}$  annimmt - deshalb nimmt man eine solche nicht an.

Die Menge der algebraischen Zahlen ist abzählbar, die Menge der transzendenten Zahlen ist überabzählbar.

## ABBILDUNG 10. Übersicht über die Teilmengen von $\mathbb{R}$

### 9. VOLLSTÄNDIGE INDUKTION

Die rekursive Definition heißt auch »Definition durch vollständige Induktion«: Um so Zahlenreihen zu definieren, wird die erste Zahl und das Gesetz zur Konstruktion einer nächsten Zahl aus ihrem Vorgänger angegeben. Beweis durch vollständige Induktion besteht also darin, eine solche Definition (Behauptung) nachzuprüfen:

- (1) Beweise: die Behauptung gilt für die erste Zahl.
- (2) Die Behauptung gelte für eine Zahl (Voraussetzung). Beweise: dann gilt die Behauptung auch für die nächste Zahl.

**Definition 40.** Summenschreibweise

**Definition.**

$$a_1 + a_2 + \dots + a_n =: \sum_{i=1}^n a_i$$

Lies: Summe über  $a_i$  von  $i = 1$  bis  $n$ . Es gilt:

(1)

$$\sum_{i=1}^n a_i = \sum_{i=1}^m a_i + \sum_{i=m+1}^n a_i$$

für  $1 \leq m \leq n$ . Bei  $n = 1$  besteht per Vereinbarung die leere Summe.

(2) Assoziativgesetz für  $n$  Summanden:

$$\lambda \cdot \sum_{i=1}^n a_i = \sum_{i=1}^n \lambda a_i$$

(3)

$$\lambda \cdot \sum_{i=1}^n a_i + \mu \cdot \sum_{j=1}^n b_j = \sum_{i=1}^n \lambda a_i + \sum_{i=1}^n \mu b_i = \sum_{i=1}^n (\lambda a_i + \mu b_i)$$

(4) Wenn die Summanden alle identisch, also unabhängig von  $i$  sind, schreibt man:

$$\sum_{i=1}^n a = n \cdot a$$

**Definition 41.** Beweis durch vollständige Induktion

**Definition.** Sei  $A(n)$  eine Aussage  $A$  über das Element  $a_n$ ,  $n \in \mathbb{Z}$  einer nummerierbaren Zahlenfolge. Schema zum Beweis von  $A(n)$  für alle  $n \geq n_0$ :

- (1) Induktionsbeginn<sup>4</sup>: Zeige  $A(n)$  für  $n = n_0$ .
- (2) Induktionsannahme:  $A(n)$  sei richtig für ein  $n \geq n_0$  (in Variation:  $A(n)$  sei richtig für alle  $n$  mit  $n_0 \leq n \leq k$ ).
- (3) Induktionsschluss<sup>5</sup> von  $n$  auf  $n+1$ : Zeige  $A(n+1)$  unter Verwendung der Induktionsannahme.  $\Rightarrow$  Die Aussage  $A(n)$  ist richtig für alle  $n \geq n_0$ ,  $n \in \mathbb{Z}$ .

**Example 25.** Beweis der Gauß-Formel durch vollständige Induktion

**Example.** Behauptung:  $\sum_{k=1}^n k = \frac{n(n+1)}{2}$  für alle  $n \in \mathbb{N}$ ,  $n \geq n_0 = 1$ .

- (1) Induktionsbeginn: Zeige  $A(n_0) = A(1)$ . Zu zeigen ist also:

$$\begin{aligned} \sum_{k=1}^1 k &\stackrel{!}{=} \frac{1(1+1)}{2} \\ \Rightarrow 1 &= 1 \quad (w) \end{aligned}$$

- (2) Induktionsannahme: Es gibt ein  $n$ , für das die Behauptung  $\sum_{k=1}^n k = \frac{n(n+1)}{2}$ .
- (3) Induktionsschluss von  $n$  auf  $n+1$ : Man überlege sich an dieser Stelle zuerst, was man zeigen will. Nämlich: Man zeige allgemein, dass aus  $A(n)$  (die aufgrund der Induktionsannahme als richtig angenommen wurde)  $A(n+1)$  folgt, also:

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

Der Beweis ist ein einfacher direkter Beweis, bei dem man aus dem linken Teil von  $A(n+1)$  unter Verwendung der Annahme  $A(n)$  den rechten Teil von  $A(n+1)$  folgert:

$$\sum_{k=1}^{n+1} k = \left( \sum_{k=1}^n k \right) + (n+1)$$

Dabei muss irgendwo die Annahme angewendet werden, nämlich hier:

$$\left( \sum_{k=1}^n k \right) + (n+1) = \frac{n(n+1)}{2} + n+1 = \frac{n(n+1) + 2(n+1)}{2} = \frac{(n+1)(n+2)}{2}$$

**Example 26.** Beweis durch vollständige Induktion von  $\sum_{k=1}^n (2k-1) = n^2$ 

**Example.** Behauptung: »Addiert man die ersten  $n$  ungeraden Zahlen, erhält man  $n^2$ «, als Formel  $\sum_{k=1}^n (2k-1) = n^2$ ,  $n_0 = 1$ .

- (1) Induktionsbeginn:

$$\begin{aligned} \sum_{k=1}^1 (2k-1) &\stackrel{!}{=} 1^2 \\ \Rightarrow 1 &= 1 \quad (w) \end{aligned}$$

- (2) Induktionsannahme: Gelte die Formel für ein beliebiges  $n \geq n_0$ , d.h. gelte  $\sum_{k=1}^n (2k-1) = n^2$ .
- (3) Induktionsschluss von  $n$  auf  $n+1$ : Gezeigt werden soll unter Verwendung der Induktionsannahme durch direkten Beweis:

$$\sum_{k=1}^{n+1} (2k-1) \stackrel{!}{=} (n+1)^2$$

Beweis:

$$\begin{aligned} \sum_{k=1}^{n+1} (2k-1) &= \left( \sum_{k=1}^n (2k-1) \right) + 2(n+1) - 1 \\ &= n^2 + 2n + 1 \\ &= (n+1)^2 \end{aligned}$$

qed.

<sup>4</sup>auch: Induktionsverankerung.

<sup>5</sup>von »Schließen auf«; auch: Induktionsschritt

**Example 27.** Beweis durch vollständige Induktion von  $\sum_{i=0}^n q^i = \frac{1-q^{n+1}}{1-q}$

**Example.** Behauptung:

$$(9) \quad \sum_{i=0}^n q^i = \frac{1-q^{n+1}}{1-q}$$

für  $q \neq 1, n \in \mathbb{Z}, n \geq 0$ .

(1) Induktionsbeginn:  $n_0 = 0$

$$\begin{aligned} \sum_{i=0}^0 q^i &\stackrel{!}{=} \frac{1-q^1}{1-q} \\ \Leftrightarrow q^0 &= 1 \quad (w) \end{aligned}$$

(2) Induktionsannahme: Für ein  $n \geq 0$  sei die Behauptung richtig.

(3) Induktionsschluss: Zu zeigen ist:  $\sum_{i=0}^{n+1} q^i \stackrel{!}{=} \frac{1-q^{n+2}}{1-q}$ . Beweis:

$$\begin{aligned} \sum_{i=0}^{n+1} q^i &= \left( \sum_{i=0}^n q^i \right) + q^{n+1} \\ &= \frac{1-q^{n+1}}{1-q} + q^{n+1} = \frac{1-q^{n+1} + q^{n+1} - q^{n+2}}{1-q} \\ &= \frac{1-q^{n+2}}{1-q} \end{aligned}$$

qed.

Folgerung: Für  $a \neq b$  gilt:

$$\frac{a^{n+1} - b^{n+1}}{a - b} = a^n + a^{n-1}b + a^{n-2}b^2 + a^{n-3}b^3 + \dots + ab^{n-1} + b^n$$

Direkter Beweis: Für  $q = \frac{b}{a}$  liefert Gleichung 9:

$$\begin{aligned} 1 + \frac{b}{a} + \left(\frac{b}{a}\right)^2 + \dots + \left(\frac{b}{a}\right)^n &= \frac{1 - \left(\frac{b}{a}\right)^{n+1}}{1 - \frac{b}{a}} \quad | \cdot a^n \\ \Leftrightarrow a^n + a^{n-1}b + a^{n-2}b^2 + a^{n-3}b^3 + \dots + ab^{n-1} + b^n &= \frac{a^n - \frac{b^{n+1}}{a}}{1 - \frac{b}{a}} \\ &= \frac{a^{n+1} - b^{n+1}}{a - b} \end{aligned}$$

**Example 28.** Beweis der Bernoulli-Ungleichung mit vollständiger Induktion

**Example.** Die Bernoulli-Ungleichung ist:  $(1+h)^n > 1+nh$  für  $h > -1, h \neq 0, n \geq 2$ .

(1) Induktionsbeginn:  $n_0 = 2$ , dafür behauptet die Bernoulli-Ungleichung:

$$\begin{aligned} (1+h)^2 &\stackrel{!}{>} 1+2h \\ \Rightarrow 1+2h+h^2 &> 1+2h \quad (w) \end{aligned}$$

denn  $h \neq 0$ .

(2) Induktionsannahme: Die Behauptung sei richtig für ein beliebiges  $n \geq 2$ , d.h. es gilt:

$$(1+h)^n > (1+nh)$$

(3) Induktionsschluss: Zu zeigen ist:

$$(1+h)^{n+1} \stackrel{!}{>} (1+nh+h)$$

unter Voraussetzung und Verwendung der Induktionsannahme.

$$\begin{aligned} (1+h)^{n+1} &= (1+h)^n (1+h) \\ &> (1+nh)(1+h) \end{aligned}$$

Hier wurde die Induktionsannahme und das Gesetz der Monotonie der Multiplikation verwendet:  $a > b \Rightarrow ac > bc$  für  $c > 0$  (vgl. »Grundgesetze der Anordnung von  $\mathbb{R}$  in diesem Skript«).

$$\begin{aligned} \Leftrightarrow (1+h)^{n+1} &> 1+nh+h+nh^2 \\ &> 1+nh+h \end{aligned}$$

qed.

Bemerkung. Für  $n = 1$  in der Bernoulli-Gleichung gilt  $1 + h = 1 + h$ . Daher gilt die Bernoulli-Gleichung auch in folgender Form:

$$(1 + h)^n \geq 1 + nh$$

für  $h > -1$ ,  $h \in \mathbb{R}$ ,  $h \neq 0$ ,  $n \geq 1$ .

## 10. BINOMIALKOEFFIZIENTEN

**Definition 42.** Binom

**Definition.** Ein Binom ist ein Ausdruck der Form  $(a + b)^n$ .

Ausrechnen von Polynomen:

$$(a + b)^0 = 1$$

$$(a + b)^1 = 1 \cdot a + 1 \cdot b$$

$$(a + b)^2 = 1 \cdot a^2 + 2 \cdot ab + 1 \cdot b^2$$

$$(a + b)^3 = 1 \cdot a^3 + 3 \cdot a^2b + 3 \cdot ab^2 + 1 \cdot b^3$$

Die Koeffizienten für Polynome beliebigen Grades können einfach mit dem Pascal'schen Dreieck berechnet werden; die Koeffizienten ergeben sich jeweils aus der Summe der darüberliegenden Koeffizienten:

$$\begin{array}{ccccccc} & & & & 1 & & \\ & & & 1 & & 1 & \\ & & 1 & & 2 & & 1 \\ & 1 & & 3 & & 3 & & 1 \\ 1 & & 4 & & 6 & & 4 & & 1 \end{array}$$

**Definition 43.** Fakultät in rekursiver Definition

**Definition.**  $0! := 1$ ,  $(n + 1)! := n!(n + 1)$  für  $n = 0, 1, 2, \dots$ . Dies ist eine rekursive Definition, d.h. eine Definition durch vollständige Induktion. Ausführlich würde man schreiben:

- (1) Definitionsbeginn:  $0! := 1$ .
- (2) Definitionsannahme: Sei  $n!$  definiert für ein  $n \geq 0$ .
- (3) Definitionsschritt von  $n$  zu  $n + 1$ : Setze  $(n + 1)! := n! \cdot (n + 1)$ .

**Definition 44.** Binomialkoeffizient

**Definition.**

$$\binom{\alpha}{k} := \frac{\alpha \cdot (\alpha - 1) \cdot (\alpha - 2) \cdot \dots \cdot (\alpha - k + 1)}{k!}$$

mit  $\alpha \in \mathbb{R}$ ,  $k \in \mathbb{N}$ . Lies: » $\alpha$  über  $k$ «. Der Zähler hat  $k$  Faktoren.

**Example 29.** Binomialkoeffizienten

**Example.**

$$\binom{6}{3} = \frac{6 \cdot 5 \cdot 4}{1 \cdot 2 \cdot 3} = 20$$

$$\binom{0,2}{4} = \frac{0,2 \cdot (0,2 - 1) \cdot (0,2 - 2) \cdot (0,2 - 3)}{1 \cdot 2 \cdot 3 \cdot 4}$$

Sei  $n \in \mathbb{N}$ . Es gilt:

$$\begin{aligned} \binom{n}{k} &= \frac{n \cdot (n - 1) \cdot (n - 2) \cdot \dots \cdot (n - k + 1)}{1 \cdot 2 \cdot 3 \cdot \dots \cdot k} \cdot \frac{(n - k) \cdot (n - k - 1) \cdot \dots \cdot 1}{(n - k) \cdot (n - k - 1) \cdot \dots \cdot 1} \\ &= \frac{n!}{k! (n - k)!} \end{aligned}$$

Aus dieser Schreibweise für Binomialkoeffizienten mit der Fakultät folgt:

$$\binom{n}{n - k} = \frac{n!}{(n - k)! \cdot (n - (n - k))!} = \frac{n!}{(n - k)! \cdot k!} = \binom{n}{k}$$

Damit ist gezeigt:

(1)

$$\binom{n}{k} = \frac{n!}{k! (n - k)!}$$

$$(10) \quad \binom{n}{k} = \binom{n}{n-k}$$

Gleichung 10 hilft zur Vereinfachung von Rechenoperationen:

$$\binom{10}{8} = \binom{10}{2} = \frac{10 \cdot 9}{1 \cdot 2} = 45$$

Spezialfälle sind:

$$\binom{n}{n} = \frac{n!}{n! \cdot 0!} = \frac{n!}{n!} = 1 =: \binom{n}{0}$$

**Definition 45.** Hilfsformel zum binomischen Lehrsatz

**Definition.**

$$\binom{n}{k} + \binom{n}{k-1} = \binom{n+1}{k}$$

für  $n, k \in \mathbb{N}$ . Beweis:

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

**Definition 46.** Binomischer Lehrsatz

**Definition.** Für  $n \geq 1$  gilt:

$$\begin{aligned} (a+b)^n &= \binom{n}{0} a^n + \binom{n}{1} a^{n-1} b + \binom{n}{2} a^{n-2} b^2 + \dots + \binom{n}{n} b^n \\ &= \sum_{k=0}^n \binom{n}{k} \cdot a^{n-k} \cdot b^k \end{aligned}$$

Beweis zu Definition 46.

(1) Induktionsbeginn:  $n_0 = 1$ . Zu zeigen:  $(a+b)^1 \stackrel{!}{=} \binom{1}{0} a^1 + \binom{1}{1} b^1$ . Beweis:

$$\begin{aligned} \binom{1}{0} a^1 + \binom{1}{1} b^1 &= a^1 + b^1 \\ &= (a+b)^1 \end{aligned}$$

qed.

(2) Induktionsannahme: Die Behauptung in Definition 46 sei richtig für ein  $n \geq 1$ .

(3) Induktionsschluss von  $n$  auf  $n+1$ : Zu zeigen ist unter Verwendung der Induktionsannahme:

$$(a+b)^{n+1} = \binom{n+1}{0} a^{n+1} + \binom{n+1}{1} a^n b + \binom{n+1}{2} a^{n-1} b^2 + \dots + \binom{n+1}{n+1} b^{n+1}$$

Beweis. Allgemein versucht man, sofern eine Gleichung zu beweisen ist, ihre linke Seite so umzuformen, dass man die Annahme einsetzen kann:

$$\begin{aligned} (a+b)^{n+1} &= (a+b)^n \cdot (a+b) \\ &= \left( \binom{n}{0} a^n + \binom{n}{1} a^{n-1} b + \binom{n}{2} a^{n-2} b^2 + \dots + \binom{n}{n-1} a b^{n-1} + \binom{n}{n} b^n \right) \cdot (a+b) \\ &= \binom{n}{0} a^{n+1} + \binom{n}{1} a^n b + \binom{n}{2} a^{n-1} b^2 + \dots + \binom{n}{n-1} a^2 b^{n-1} + \binom{n}{n} a b^n \\ &\quad + \binom{n}{0} a^n b + \binom{n}{1} a^{n-1} b^2 + \binom{n}{2} a^{n-2} b^3 + \dots + \binom{n}{n-1} a b^n + \binom{n}{n} b^n \end{aligned}$$



Unter Verwendung von Definition 45 lässt sich diese Gleichung wie folgt zusammenfassen:

$$\begin{aligned}(a+b)^{n+1} &= \binom{n+1}{0} a^{n+1} + \binom{n+1}{1} a^n b + \binom{n+1}{2} a^{n-1} b^2 + \dots \\ &\quad + \binom{n+1}{n-1} a^2 b^{n-1} + \binom{n+1}{n} a b^n + \binom{n+1}{n+1} b^{n+1}\end{aligned}$$

qed. Auf dieser Formel beruht das Pascalsche Dreieck!

Warum ist die Konstruktion des Pascalschen Dreiecks richtig? Der Binomische Lehrsatz besagt:

$$(a+b)^n = \binom{n}{0} a^n + \binom{n}{1} a^{n-1} b + \binom{n}{2} a^{n-2} b^2 + \dots + \binom{n}{n} b^n$$

Am Anfang ergänzt man einen weiteren Koeffizienten:

$$\binom{n+1}{0}$$

Der zweite Koeffizient ist die Summe des ersten und zweiten Koeffizienten ist:

$$\binom{n}{0} + \binom{n}{1} = \binom{n+1}{1}$$

Der dritte Koeffizient ist die Summe des zweiten und dritten Koeffizienten ist:

$$\binom{n}{1} + \binom{n}{2} = \binom{n+1}{2}$$

**Example 30.** Berechnung von  $(a+b)^5$  mit dem Binomischen Lehrsatz

**Example.** Nach obiger Gleichung gilt folgendes:

$$\begin{aligned}(a+b)^5 &= \sum_{k=0}^5 \binom{5}{k} \cdot a^{5-k} \cdot b^k \\ &= \binom{5}{0} a^5 b^0 + \binom{5}{1} a^4 b^1 + \binom{5}{2} a^3 b^2 + \binom{5}{3} a^2 b^3 + \binom{5}{4} a^1 b^4 + \binom{5}{5} a^0 b^5 \\ &= a^5 + 5a^4 b + 10a^3 b^2 + 10a^2 b^3 + 5a^1 b^4 + b^5\end{aligned}$$

**Example 31.** Berechnung von  $(a-b)^5$  mit dem Binomischen Lehrsatz

**Example.** Es muss nun darauf geachtet werden, die Vorzeichen zwischen den einzelnen Summanden richtig zu setzen. Wie bereits bekannt gilt:

$$\begin{aligned}(a-b)^5 &= \sum_{k=0}^5 \binom{5}{k} \cdot a^{5-k} \cdot b^k \\ &= \binom{5}{0} a^5 b^0 - \binom{5}{1} a^4 b^1 + \binom{5}{2} a^3 b^2 - \binom{5}{3} a^2 b^3 + \binom{5}{4} a^1 b^4 - \binom{5}{5} a^0 b^5 \\ &= a^5 - 5a^4 b + 10a^3 b^2 - 10a^2 b^3 + 5a^1 b^4 - b^5\end{aligned}$$

## 11. WINKELFUNKTIONEN

11.1. **Winkel.** Was sind Winkel und wie misst man sie? Indem man den Vollwinkel gleich einer beliebigen Zahl  $k$  setzt und Winkel als Bruchteile von  $k$  angibt. Dazu sind verschiedene Systeme in Gebrauch:

**Einheitssystem::** Man setzt  $k = 1$

**Babylonisches System::**  $k = 360^\circ$  (»Grad«)

**Neugrad::**  $k = 400$  (»Neugrad«)

**Bogenmaß::** Einzig sinnvoll ist die Wahl  $k = 2\pi$  (»Bogenmaß« bzw. »Radiant«), denn  $2\pi$  entspricht dem Umfang eines Kreises mit dem Radius 1. In der Mathematik wird allein mit Bogenmaß gerechnet. Abschätzung von Bogenmaßwerten:  $1,5 \approx 90^\circ$ ,  $6 \approx 360^\circ$ . Vorsicht: am Taschenrechner muss die Winkleinheit stets richtig eingestellt werden. Das Bogenmaß als Winkelmaß gibt gleichzeitig die Länge des Bogens  $\alpha$  zu einem Winkel  $\alpha$  an, wenn ein Kreis mit Radius  $r = 1$  um das Zentrum des Winkels betrachtet wird.

Positive Winkel werden in der Mathematik stets im mathematischen Drehsinn (d.i. gegen den Uhrzeigersinn) gemessen; negative Winkel sind solche, die mit dem Uhrzeigersinn gemessen werden.

Winkel  $\alpha > 2\pi$  geben Winkel bei mehrfacher Umdrehung im Kreis an man erkennt daran, wie oft der volle Kreisumfang zurückgelegt wurde, um zu der aktuellen Position zu kommen.

11.2. **Einführung.** Zwei Dreiecke heißen ähnlich, wenn sie in zwei Winkeln übereinstimmen. Der einzige Satz über ähnliche Dreiecke (»Hauptsatz«) besagt: entsprechende Dreiecksseiten stehen im gleichen Verhältnis. Also:

$$\frac{b}{c} = \frac{b'}{c'}$$

$$\frac{b}{a} = \frac{b'}{a'}$$

Angewandt auf rechtwinklige Dreiecke heißt das: rechtwinklige Dreiecke sind ähnlich, wenn sie in einem Winkel  $\phi$  außer dem rechten Winkel übereinstimmen. Dann stehen wiederum entsprechende Seiten im gleichen Verhältnis. Unter Verwendung der Seitenbezeichnungen Ankathete  $A$ , Gegenkathete  $G$  und Hypotenuse  $H$  heißt das:

$$\frac{G}{A} = \frac{G'}{A'} = \tan \phi$$

$$\frac{A}{H} = \frac{A'}{H'} = \cos \phi$$

$$\frac{G}{H} = \frac{G'}{H'} = \sin \phi$$

Diese Verhältnisse hängen nur von einem Winkel ab und werden deshalb als Winkelfunktionen bezeichnet. Definition der Winkelfunktionen unabhängig von Winkeln und Winkleinheiten (siehe Abbildung Abb.Einheitskreis.eps):

- Der Cosinus ist die Projektion des Radius des Einheitskreises auf die  $x$ -Achse.
- Der Sinus ist die Projektion des Radius des Einheitskreises auf die  $y$ -Achse.
- Der Tangens ist die Projektion des bis zu einer senkrechten Tangente verlängerten Radius auf die  $y$ -Achse. Deshalb ist  $\lim_{\phi \rightarrow \frac{\pi}{2}} \tan \phi = \infty$ , d.h. die Tangenskurve ist für  $\phi = (2n+1)\frac{\pi}{2}$ ,  $n \in \mathbb{Z}$  (das ist für ungerade Vielfache von  $\frac{\pi}{2}$ ) nicht definiert. Der Tangens ist im Intervall  $]-\frac{\pi}{2}; \frac{\pi}{2}[$  eine bijektive Abbildung, so dass dazu die Umkehrfunktion Arkustangens  $\arctan$  gebildet werden kann, die jeder reellen Zahl einen Winkel  $]-\frac{\pi}{2}; \frac{\pi}{2}[$  zuordnet.

Aus der Rotation des Radius des Einheitskreises ergeben sich Sinus- und Cosinuskurve als Auftragung der Sinus- bzw. Cosinuswerte gegen den Winkel im Bogenmaß. Unter Linux können die Auftragungen der Kurven leicht durch `plot sin(x), cos(x)` im Programm `gnuplot` erstellt werden.

Additionstheoreme der Winkelfunktionen.

$$\begin{aligned}\sin(\alpha + \beta) &= \sin \alpha \cdot \cos \beta + \cos \alpha \cdot \sin \beta \\ \cos(\alpha + \beta) &= \cos \alpha \cdot \cos \beta - \sin \alpha \cdot \sin \beta \\ \tan(\alpha + \beta) &= \frac{\tan \alpha + \tan \beta}{1 - \tan \alpha \cdot \tan \beta}\end{aligned}$$

Beweis der Additionstheoreme:

- Zuerst soll gezeigt werden, dass die Gleichung 14 richtig ist, wenn die Gleichungen ?? und 13 richtig sind (diese werden in folgendem Rechenweg angewandt):

$$\begin{aligned}\tan &= \frac{\sin}{\cos} \\ \Rightarrow \tan(\alpha + \beta) &= \frac{\sin(\alpha + \beta)}{\cos(\alpha + \beta)} \\ &= \frac{\sin \alpha \cdot \cos \beta + \cos \alpha \cdot \sin \beta}{\cos \alpha \cdot \cos \beta - \sin \alpha \cdot \sin \beta} \\ &= \frac{\frac{\sin \alpha \cdot \cos \beta}{\cos \alpha \cdot \cos \beta} + \frac{\cos \alpha \cdot \sin \beta}{\cos \alpha \cdot \cos \beta}}{\frac{\cos \alpha \cdot \cos \beta}{\cos \alpha \cdot \cos \beta} - \frac{\sin \alpha \cdot \sin \beta}{\cos \alpha \cdot \cos \beta}} \\ &= \frac{\frac{\sin \alpha}{\cos \alpha} + \frac{\sin \beta}{\cos \beta}}{\frac{\cos \beta}{\cos \beta} - \frac{\sin \alpha}{\cos \alpha} \cdot \frac{\sin \beta}{\cos \beta}} \\ &= \frac{\tan \alpha + \tan \beta}{1 - \tan \alpha \cdot \tan \beta}\end{aligned}$$

- Nun soll gezeigt werden, dass die Gleichung 13 richtig ist, wenn die Gleichung ?? richtig ist. Dazu werden folgende Beziehungen verwendet, die unmittelbar aus der Darstellung der Winkelfunktionen

am Einheitskreis abgeleitet werden können:

$$\begin{aligned}\cos\left(\frac{\pi}{2} - \alpha\right) &= \sin \alpha \\ \sin\left(\frac{\pi}{2} - \alpha\right) &= \cos \alpha \\ \cos(-\alpha) &= \cos \alpha \\ \sin(-\alpha) &= -\sin \alpha\end{aligned}$$

Und nun der Rechenweg:

$$\begin{aligned}\cos(\alpha + \beta) &= \sin\left(\frac{\pi}{2} - (\alpha + \beta)\right) \\ &= \sin\left(\left(\frac{\pi}{2} - \alpha\right) + (-\beta)\right) \\ &= \sin\left(\frac{\pi}{2} - \alpha\right) \cdot \cos(-\beta) + \cos\left(\frac{\pi}{2} - \alpha\right) \cdot \sin(-\beta) \\ &= \cos \alpha \cdot \cos \beta - \sin \alpha \cdot \sin \beta\end{aligned}$$

- Also muss nur die Gleichung ?? bewiesen werden; die anderen beiden Additionstheoreme basieren darauf, wie oben gezeigt werden konnte, sind dann also ebenfalls richtig. Vergleiche zu folgendem Beweis Abbildung Abb.AddTheoremSin.eps.

$$\begin{aligned}\sin(\alpha + \beta) &= CB = CG + GB \\ \frac{CG}{CE} &= \cos \alpha \\ CG &= \cos \alpha \cdot CE \\ &= \cos \alpha \cdot \sin \beta \\ \frac{GB}{FH} &= \sin \alpha \\ \frac{GB}{FA} &= \sin \alpha \cdot FA \\ &= \sin \alpha \cdot \cos \beta\end{aligned}$$

So dass die zu beweisende Formel folgt:

$$\sin(\alpha + \beta) = \sin \alpha \cdot \cos \beta + \cos \alpha \cdot \sin \beta$$

Es gilt für jeden Winkel (herzuleiten aus den ähnlichen Dreiecken am Einheitskreis):

$$\frac{\sin \alpha}{\cos \alpha} = \frac{\tan \alpha}{1}$$

Trigonometrischer Satz des Pythagoras. Durch Anwendung des Satzes des Pythagoras auf die am Einheitskreis dargestellten Winkelfunktionen folgt der sog. trigonometrische Satz des Pythagoras:

$$\sin^2 \alpha + \cos^2 \alpha = 1$$

Damit ist nach  $\cos \alpha = \sqrt{1 - \sin^2 \alpha}$  eine Umrechnung zwischen Sinus- und Cosinuswerten möglich, jedoch muss wiederum unterschieden werden, in welchem Quadranten der Winkel  $\alpha$  liegt.

**Definition 47.** Winkelfunktionen am rechtwinkligen Dreieck

**Definition.** Sei ein Dreieck entsprechend Abbildung 11. Dann sind die Winkelfunktionen wie folgt definiert:

$$\begin{aligned}\sin(\alpha) &= \frac{a}{c} \\ \cos(\alpha) &= \frac{b}{c} \\ \tan(\alpha) &= \frac{a}{b} \\ \cot(\alpha) &= \frac{b}{a} = \frac{1}{\tan(\alpha)}\end{aligned}$$

Folgerungen aus diesen Definitionen:

ABBILDUNG 11. Übliche Bezeichnungen am Dreieck

$$\begin{aligned}\tan(\alpha) &= \frac{\sin(\alpha)}{\cos(\alpha)} \\ \tan(\alpha) \cdot \cot(\alpha) &= 1\end{aligned}$$

$$(11) \quad \sin^2(\alpha) + \cos^2(\alpha) = 1$$

Gleichung 11 ist der sog. »trigonometrische Satz des Pythagoras«. Die Winkelfunktionen haben folgende spezielle Funktionswerte:

|     | $0^\circ \triangleq 0$ | $30^\circ \triangleq \frac{\pi}{6}$ | $45^\circ \triangleq \frac{\pi}{4}$ | $60^\circ \triangleq \frac{\pi}{3}$ | $90^\circ \triangleq \frac{\pi}{2}$ |
|-----|------------------------|-------------------------------------|-------------------------------------|-------------------------------------|-------------------------------------|
| sin | 0                      | $\frac{1}{2}$                       | $\frac{\sqrt{2}}{2}$                | $\frac{\sqrt{3}}{2}$                | 1                                   |
| cos | 1                      | $\frac{\sqrt{3}}{2}$                | $\frac{\sqrt{2}}{2}$                | $\frac{1}{2}$                       | 0                                   |
| tan | 0                      | $\frac{\sqrt{3}}{3}$                | 1                                   | $\sqrt{3}$                          | $\pm\infty$                         |
| cot | $\pm\infty$            | $\sqrt{3}$                          | 1                                   | $\frac{\sqrt{3}}{3}$                | 0                                   |

**Definition 48.** Thaleskreis

**Definition.** Ein Halbkreis. Alle eingezeichneten Dreiecke, die den Durchmesser als eine Seite haben, sind rechtwinklige Dreiecke. Der Verlauf des Sinus lässt sich hier gut beobachten.

**Definition 49.** Sinussatz

**Definition.**

$$\frac{\sin(\alpha)}{\sin(\beta)} = \frac{a}{b}$$

**Definition 50.** Cosinussatz

**Definition.**

$$a^2 = b^2 + c^2 - 2bc \cdot \cos(\alpha)$$

**Definition 51.** Flächensatz

**Definition.**

$$F = \frac{1}{2}bc \cdot \sin(\alpha)$$

## ABBILDUNG 12. Winkelfunktionen am Einheitskreis

11.3. **Definition der Winkelfunktionen für beliebige Winkel.** Bezogen auf Abbildung 12 ist:

$$\sin(\alpha) = \frac{v}{r}$$

$$\cos(\alpha) = \frac{u}{r}$$

$$\tan(\alpha) = \frac{v}{u}$$

$$\cot(\alpha) = \frac{u}{v}$$

Die Koordinaten des Punktes  $P(U; V)$  beim Umlauf auf dem Kreis sind nun die Werte für  $u$  und  $v$  bei beliebigem Winkel  $\alpha \in \mathbb{R}$ .  $u$  und  $v$  können also auch negativ sein ( $r$  ist dagegen immer positiv)! Beachte: Funktionswerte für beliebige Winkel sind durch die Funktionswerte für Winkel zwischen  $0^\circ$  und  $90^\circ$  darstellbar, aufgrund der Symmetrie der Winkelfunktionen. Nach Abbildung 13 ergibt sich so für Sinuswerte für Winkel

## ABBILDUNG 13. Achsensymmetrie der Sinusfunktion

größer  $90^\circ$ :

$$\sin(150^\circ) = \frac{v}{r} = \sin(30^\circ)$$

Nach Abbildung 13 ergibt sich so für Sinuswerte für Winkel größer  $180^\circ$ :

## ABBILDUNG 14. Punktsymmetrie der Sinusfunktion

$$\sin(205^\circ) = \frac{v}{r} = -\sin(25^\circ)$$

**Example 32.** Rückführung von Winkelfunktionen bei Winkeln größer  $90^\circ$

**Example.**

$$\cos(150^\circ) = -\cos(30^\circ)$$

$$\tan(105^\circ) = -\tan(75^\circ)$$

$$\cot(290^\circ) = -\cot(70^\circ)$$

Der Winkel  $\alpha'$ , auf den man den Wert einer Winkelfunktion von  $\alpha$  zurückführen kann, ist der Ergänzungswinkel von  $\alpha$  zum nächsten Vielfachen von  $90^\circ$ .

11.4. **Bogenmaß.** Gegeben sei ein Einheitskreis, d.h. ein Kreis mit Radius  $r = 1$  (vgl. Abbildung 15). Der

ABBILDUNG 15. Der Einheitskreis

Umfang dieses Kreises ist  $U = 2\pi r = 2\pi$ . Mit einem Winkel  $\alpha$  wird vom Umfang ein Bogen  $b$  ausgeschnitten. Für dessen Läng gilt:

$$\begin{aligned} \frac{\alpha}{360^\circ} &= \frac{b}{2\pi} \\ \Rightarrow b &= \frac{\alpha \cdot 2\pi}{360} = \frac{\alpha\pi}{180} \end{aligned}$$

Diese Länge des Bogens  $b$  heißt Bogenmaß zum Winkel  $\alpha$  und ist das hier in Zukunft verwandte Maß für den Winkel. Die Einheit ist:

$$[b] = 1\text{rad} = 1\text{Radian}$$

Das Bogenmaß wird als Vielfaches von  $\pi = 3,141592654$  angegeben:

| Grad | $0^\circ$ | $30^\circ$      | $45^\circ$      | $60^\circ$      | $90^\circ$      | $180^\circ$ | $270^\circ$      | $360^\circ$ |
|------|-----------|-----------------|-----------------|-----------------|-----------------|-------------|------------------|-------------|
| rad  | 0         | $\frac{\pi}{6}$ | $\frac{\pi}{4}$ | $\frac{\pi}{3}$ | $\frac{\pi}{2}$ | $\pi$       | $\frac{3\pi}{2}$ | $2\pi$      |

Es gibt auch negative Winkel im Bogenmaß, d.h. die Winkelfunktionen haben einen Definitionsbereich von  $\mathbb{R}$ .

11.5. **Graphen der Winkelfunktionen.** Die Werte aller Winkelfunktionen eines Winkels  $\alpha$  können unmittelbar am Einheitskreis abgetragen werden:

ABBILDUNG 16. Abtragung von Sinuswerten am Einheitskreis

- Die Periode beider Graphen ist  $2\pi$ :

$$\sin x = \sin(x + k \cdot 2\pi), \quad k \in \mathbb{Z}$$

$$\cos x = \cos(x + k \cdot 2\pi), \quad k \in \mathbb{Z}$$

- Die Graphen von Sinus und Cosinus entstehen auseinander durch Verschiebung um  $\frac{\pi}{2}$ :

$$\sin x = \cos\left(x - \frac{\pi}{2}\right)$$

$$\cos x = \sin\left(x + \frac{\pi}{2}\right)$$

- Symmetrie von Cosinus und Sinus:
  - Cosinus ist achsensymmetrisch zur  $y$ -Achse:

$$\cos x = \cos(-x)$$

- Sinus ist punktsymmetrisch zum Koordinatenursprung:

$$\sin x = -\sin(-x)$$

**Definition 52.** gerade und ungerade Funktion

**Definition.** Eine Funktion  $f(x)$  heißt gerade, wenn sie symmetrisch zur  $y$ -Achse ist, d.h. wenn gilt:

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

Eine Funktion  $f(x)$  heißt ungerade, wenn sie punktsymmetrisch zum Koordinatenursprung ist, d.h. wenn gilt:

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

## ABBILDUNG 17. Konstruktion der Tangensfunktion am Einheitskreis

Der Tangens hat seinen Namen daher, dass beim Abtragen am Einheitskreis eine Tangente zur Konstruktion verwendet wird (vgl. Abbildung 17). Zur Konstruktion des Cotangens am Einheitskreis verwendet man eine andere Tangente (die »Kotangente«), die senkrecht zur Tangente steht, die zur Konstruktion des Tangens verwendet wurde.

- Die Periode von  $\tan x$ ,  $\cot x$  ist  $\pi$ :

$$\tan x = \tan(x + k\pi)$$

$$\cot x = \cot(x + k\pi)$$

- $\tan x$  und  $\cot x$  können nicht durch Verschiebung ineinander überführt werden, sondern es gilt:

$$\tan x = \frac{1}{\cot x}$$

- $\tan x$  und  $\cot x$  sind beide punktsymmetrisch zum Koordinatenursprung

**11.6. Additionstheoreme.** Die folgenden beiden Additionstheoreme werden hier ohne Beweis eingeführt, alle anderen sind daraus herzuleiten:

$$(12) \quad \sin(\alpha \pm \beta) = \sin \alpha \cdot \cos \beta \pm \cos \alpha \cdot \sin \beta$$

$$(13) \quad \cos(\alpha \pm \beta) = \cos \alpha \cdot \cos \beta \mp \sin \alpha \cdot \sin \beta$$

Additionstheorem für Tangens:

$$(14) \quad \tan(\alpha \pm \beta) = \frac{\tan \alpha \pm \tan \beta}{1 \mp \tan \alpha \cdot \tan \beta}$$

Beweis zu Gleichung 14.

$$\begin{aligned} \tan(\alpha + \beta) &= \frac{\sin(\alpha + \beta)}{\cos(\alpha + \beta)} \\ &= \frac{\sin \alpha \cdot \cos \beta + \cos \alpha \cdot \sin \beta}{\cos \alpha \cdot \cos \beta - \sin \alpha \cdot \sin \beta} \cdot \frac{\frac{1}{\cos \alpha \cdot \cos \beta}}{\frac{1}{\cos \alpha \cdot \cos \beta}} \\ &= \frac{\frac{\sin \alpha}{\cos \alpha} + \frac{\sin \beta}{\cos \beta}}{1 - \frac{\sin \alpha \cdot \sin \beta}{\cos \alpha \cdot \cos \beta}} = \frac{\tan \alpha + \tan \beta}{1 - \tan \alpha \cdot \tan \beta} \end{aligned}$$

Weitere Additionstheoreme. Sog. »Halbwinkelsätze«:

$$\sin \gamma + \sin \delta = 2 \cdot \sin \frac{\gamma + \delta}{2} \cdot \cos \frac{\gamma - \delta}{2}$$

$$\sin \gamma - \sin \delta = 2 \cdot \cos \frac{\gamma + \delta}{2} \cdot \sin \frac{\gamma - \delta}{2}$$

$$\cos \gamma + \cos \delta = 2 \cdot \cos \frac{\gamma + \delta}{2} \cdot \cos \frac{\gamma - \delta}{2}$$

$$\cos \gamma - \cos \delta = -2 \cdot \sin \frac{\gamma + \delta}{2} \cdot \sin \frac{\gamma - \delta}{2}$$

Setze in Gleichungen 12 und 13  $\alpha = \beta$ . Es entsteht, unter Verwendung von Gleichung 11:

$$\sin 2\alpha = 2 \sin \alpha \cdot \cos \alpha$$

$$\cos 2\alpha = \cos^2 \alpha - \sin^2 \alpha = \begin{cases} 1 - 2 \sin^2 \alpha \\ -1 + 2 \cos^2 \alpha \end{cases}$$

## 12. UMKEHRFUNKTION

## 12.1. Definition.

**Definition 53.** Umkehrfunktion

**Definition.** Seien  $D, B$  Mengen. Sei  $f : x \mapsto y$  mit  $x \in D, y \in B$  eine Bijektion. Dann heißt  $f^{-1} : y \mapsto x$  mit  $y \in B, x \in D$  und  $y = f(x)$  die Umkehrfunktion<sup>6</sup> (auch: »inverse Funktion«) zu  $f$ .

Die Bezeichnung der unabhängigen und abhängigen Variablen einer Funktion ist beliebig. Insbesondere darf man in  $f^{-1} x$  in  $y$  und  $y$  in  $x$  umbenennen! Das wird bei der Erzeugung der Umkehrfunktion angewandt!

<sup>6</sup>auch geschrieben als  $\bar{f}(x)$

**12.2. Wie erhält man  $f^{-1}$  aus  $f$ ?** Jede streng monotone Funktion  $f$  besitzt eine Umkehrfunktion  $f^{-1}(x)$ . Eine Umkehrfunktion kann nur existieren, wenn die Funktion  $f(x)$  streng monoton steigend oder fallend verläuft. Um nachzuweisen, dass  $f(x)$  streng monoton steigend (SMS) oder streng monoton fallend (SMF) ist, müssen folgende Bedingungen gelten:

- $f'(x) > 0$  im betrachteten Intervall damit SMS gilt
- $f'(x) < 0$  im betrachteten Intervall damit SMF gilt

Es kann verschiedene Intervalle geben, in denen die Monotonität unterschiedlich ist. Diese Intervalle ermittelt man, indem man die Nullstellen der ersten Ableitung ermittelt (dort gilt ja nicht mehr  $f'(x) < 0$  oder  $f'(x) > 0$ ). An diesen Nullstellen existiert keine Umkehrfunktion. Die links und rechts liegenden Intervalle müssen gesondert untersucht werden.

Verfahren zur Ermittlung der Funktionsgleichung:

- (1)  $y = f(x)$  nach  $x$  auflösen, sofern möglich. Ist dies nicht möglich, so kann die Umkehrfunktion nur rechnerisch ermittelt werden.
- (2) Benenne  $x$  in  $y$  und  $y$  in  $x$  um. Das bedeutet in einem Koordinatensystem die Spiegelung von  $f(x)$  an der Winkelhalbierenden  $w : y = x$  durch den 1. und 3. Quadranten; d.h., man erhält den Graphen von  $f^{-1}$  indem man den Graphen von  $f$  an  $w$  spiegelt.
- (3) Auch Bild- und Urbildbereich werden vertauscht.

**Example 33.** Funktion aus der Umkehrfunktion erhalten

**Example.**  $f : \mathbb{R} \mapsto \mathbb{R}$  mit  $y = 2x - 1$ . Vgl. Abbildung 18.

- (1) Auflösen nach  $x$ :

$$\begin{aligned} y &= 2x - 1 \\ \Leftrightarrow 2x &= y + 1 \\ \Leftrightarrow x &= \frac{1}{2}y + \frac{1}{2} \end{aligned}$$

- (2) Vertausche  $x$  und  $y$ . Es entsteht  $f^{-1} : \mathbb{R} \mapsto \mathbb{R}$  mit

$$y = \frac{1}{2}x + \frac{1}{2}$$

ABBILDUNG 18. Funktion und Umkehrfunktion

**12.3. Ableitung der Umkehrfunktion.** Für die Ableitung von Umkehrfunktionen gilt  $(f^{-1}(x))' = \frac{1}{f'(x)}$ . Wenn die Ableitung der ursprünglichen Funktion also bekannt ist, lässt sich leicht die Ableitung der Umkehrfunktion berechnen. Die Ableitung kann man natürlich auch auf dem traditionellen Weg berechnen: Man berechnet zuerst  $f^{-1}(x)$  und leitet diese dann ab.

### 13. DIE ARCUSFUNKTIONEN

Von lat. »arcus«: der Bogen. So genannt, weil man mit diesen Umkehrfunktionen der trigonometrischen Funktionen aus einem Funktionswert die Länge des Bogens (d.h. ein Winkelmaß im Bogenmaß) erhält. Weil die trigonometrischen Funktionen mit  $\mathbb{D} = \mathbb{R}$  periodisch und somit keine Bijektionen sind, können sie nur unter Einschränkung des Definitionsbereichs umgekehrt werden. Die Auflösung von z.B.  $y = \sin x$  zu  $x$  ist rechnerisch nicht möglich, weshalb man für die Umkehrfunktion einfach einen neuen Funktionsnamen einführt.

**13.1. Arcussinus.** Funktion:  $f : [-\frac{\pi}{2}; \frac{\pi}{2}] \mapsto [-1; 1]$  mit  $y = \sin x$ .

Umkehrfunktion:  $f^{-1} : [-1; 1] \mapsto [-\frac{\pi}{2}; \frac{\pi}{2}]$  mit  $y = \arcsin x$ , auch geschrieben als  $y = \sin^{-1} x$ .

Der Graph von  $\arcsin x$  entsteht durch Spiegelung des Graphen von  $\sin x$  an der 1. Winkelhalbierenden (vgl. Abbildung 19).

ABBILDUNG 19. Die Funktionen  $\sin x$  und  $\arcsin x$

**13.2. Arcuscosinus.** Funktion:  $f : [0; \pi] \mapsto [-1; 1]$  mit  $y = \cos x$ .

Umkehrfunktion:  $f^{-1} : [-1; 1] \mapsto [0; \pi]$  mit  $y = \arccos x$ , auch geschrieben als  $y = \cos^{-1} x$ .

Vergleiche Abbildung 20.

**13.3. Arcustangens.** Funktion:  $f : \left(-\frac{\pi}{2}; \frac{\pi}{2}\right) \mapsto (-\infty; \infty)$  mit  $y = \tan x$ .

Umkehrfunktion:  $f : (-\infty; \infty) \mapsto \left(-\frac{\pi}{2}; \frac{\pi}{2}\right)$  mit  $y = \arctan x$ .

Vergleiche Abbildung 21.  $\arctan x$  ist eine beschränkte Funktion, d.h. eine Funktion, deren Funktionswerte sich trotz unendlichem Definitionsbereich vollständig im Endlichen bewegen.

**13.4. Arcuscotangens.** Funktion:  $f : (0; \pi) \mapsto (-\infty; \infty)$  mit  $y = \cot x$ .

Umkehrfunktion:  $f : (-\infty; \infty) \mapsto (0; \pi)$  mit  $y = \operatorname{arccot} x$ .

Vergleiche Abbildung 22.

## 14. KOMPLEXE ZAHLEN

Komplexe Zahlen sind u.a. notwendig zur Fourier-Transformation.

**14.1. Definition.** Die Menge der komplexen Zahlen  $\mathbb{C}$  ist eine Obermenge von  $\mathbb{R}$ :  $\mathbb{R} \subset \mathbb{C}$ . Die komplexen Zahlen sind erforderlich, weil in  $\mathbb{R}$  manche quadratischen Gleichungen nicht lösbar sind, z.B.  $x^2 + 1 = 0$ . Man führte deshalb eine »formale« Lösung ein:

$$\begin{aligned} x^2 &= -1 \\ \Rightarrow x_{1|2} &= \pm\sqrt{-1} \end{aligned}$$

Wurzeln sind immer positiv - meint man eine negative Zahl, muss man also explizit ein negatives Vorzeichen vor die Wurzel schreiben:  $\sqrt{4} = 2$ ;  $-\sqrt{4} = -2$ ;  $\sqrt{(-2)^2} = 2$ . Zur Lösung der in  $\mathbb{R}$  nicht lösbaren quadratischen Gleichungen ist nur das Symbol  $i = \sqrt{-1}$  zusätzlich zu  $\mathbb{R}$  nötig.

**Definition 54.** imaginäre Einheit

**Definition.**  $i := \sqrt{-1}$  heißt imaginäre Einheit.

**Definition 55.** imaginäre Zahl

**Definition.**  $b \cdot i = i \cdot b$  mit  $b \in \mathbb{R}$ ,  $b \neq 0$  heißt komplexe Zahl.

**Definition 56.** komplexe Zahl

**Definition.**  $(a + bi)$  mit  $a, b \in \mathbb{R}$  heißt komplexe Zahl. Für die Komplexe Zahl  $0 + 0i$  schreibt man einfach 0.

**Definition 57.** Menge der komplexen Zahlen

**Definition.**  $\mathbb{C} := \{a + bi \mid a, b \in \mathbb{R}\}$  heißt Menge der komplexen Zahlen.

**14.2. Addition und Multiplikation komplexer Zahlen.** Man überträgt die Rechenregeln für Addition, Subtraktion, Multiplikation und Division von  $\mathbb{R}$  auf  $\mathbb{C}$  in naheliegender Weise.

$$\begin{aligned} i^2 &= -1 \\ i^3 &= i^2 \cdot i = -i \\ i^4 &= i^2 \cdot i^2 = -1 \cdot -1 = 1 \\ i^5 &= i^4 \cdot i = 1 \cdot i = i \\ i^{28} &= (i^4)^7 = 1^7 = 1 \\ i^{101} &= i \\ i^{1031} &= i^{1028} \cdot i^3 = i^3 = -i \end{aligned}$$

Addition, Subtraktion und Multiplikation werden analog zu  $\mathbb{R}$  durchgeführt, so als sei  $i$  eine Einheit:

$$\begin{aligned} (-2 + 4i) + (0,5 - 2i) &= (-2 + 0,5) + (4i - 2i) = -1,5 + 2i \\ (-3 + 5i)(2 - 4i) &= -6 + 12i + 10i - 20i^2 = 14 + 22i \end{aligned}$$

Es gelten bezüglich Addition und Multiplikation die gleichen Grundgesetze (VANIK) wie in  $\mathbb{R}$ . Achtung: In  $\mathbb{C}$  gibt es keine Anordnung, d.h. kein  $\gg$  und kein  $\ll$ . Bei Vergleichen gilt nur  $\gg$  oder  $\ll$ . Es gilt weder  $-3 + 5i < 4 - 2i$  noch  $-3 + 5i > 4 - 2i$ !

**Definition 58.** Real- und Imaginärteil

ABBILDUNG 20. Die Funktionen  $\cos x$  und  $\arccos x$

ABBILDUNG 21. Die Funktionen  $\tan x$  und  $\arctan x$

ABBILDUNG 22. Die Funktionen  $\cot x$  und  $\operatorname{arccot} x$



**Definition.** Sei eine komplexe Zahl  $z = a + bi \in \mathbb{C}$ . Dann heit:

- $a$  der Realteil von  $z$ , in Zeichen  $\operatorname{Re}(z) = a$
- $b$  der Imaginrteil von  $z$ , in Zeichen  $\operatorname{Im}(z) = b$
- $bi$  der imaginre Anteil von  $z$ .
- $\bar{z} := a - bi$  die konjugiert komplexe Zahl zu  $z$ .

**Example 34.** konjugiert komplexe Zahl

**Example.** Man bilde die konjugiert komplexen Zahlen:

- $z_1 = 3 + 2i$ ,  $\bar{z}_1 = 3 - 2i$
- $z_2 = -5 - 6i$ ,  $\bar{z}_2 = -5 + 6i$
- $z_3 = 2$ ,  $\bar{z}_3 = 2$

**Definition 59.** Wann ist eine komplexe Zahl reell?

**Definition.** Sei  $z \in \mathbb{C}$ . Es gilt  $z \in \mathbb{R} \Leftrightarrow z = \bar{z}$ . Beweis:

- Teilbeweis  $z \in \mathbb{R} \Rightarrow z = \bar{z}$ :

$$z \in \mathbb{R} \Rightarrow z = a = a + 0i = a - 0i = a = \bar{z}$$

- Teilbeweis  $z \in \mathbb{R} \Leftarrow z = \bar{z}$ . Es ist  $z = a + bi$ . Sei  $z = \bar{z}$ .

$$\begin{aligned} z &= \bar{z} \\ \Rightarrow a + bi &= a - bi \\ \Rightarrow 2bi &= 0 \\ \Rightarrow b &= 0 \end{aligned}$$

Sei  $z = a + bi$ . Dann ist:

$$z\bar{z} = (a + bi)(a - bi) = a^2 - (bi)^2 = a^2 + b^2$$

Also ist  $z\bar{z} \geq 0$  fr alle  $z$ !

**Definition 60.** Betrag einer komplexen Zahl

**Definition.** Sei  $z = a + bi$ . Dann heit  $\sqrt{a^2 + b^2}$  der Betrag von  $z$ , in Zeichen  $|z| := \sqrt{a^2 + b^2} = \sqrt{z\bar{z}}$ .

**Example 35.** Betrag einer komplexen Zahl

**Example.**  $z = -3 + 2i$ . Dann ist  $|z| = \sqrt{9 + 4} = \sqrt{13}$ .

### 14.3. Division komplexer Zahlen.

**Definition 61.** Division komplexer Zahlen

**Definition.**

$$\frac{z_1}{z_2} = \frac{z_1 \cdot \bar{z}_2}{z_2 \cdot \bar{z}_2} = \frac{z_1 \cdot \bar{z}_2}{|z_2|^2}$$

**Example 36.** Division komplexer Zahlen

**Example.**  $z_1 = 2 + 3i$ ,  $z_2 = 1 + 4i$ . Man erweitert mit dem konjugiert komplexen Nenner, um den Nenner reell zu machen:

$$\frac{z_1}{z_2} = \frac{2 + 3i}{1 + 4i} \cdot \frac{1 - 4i}{1 - 4i} = \frac{2 - 8i + 3i - 12i^2}{1 + 16} = \frac{14}{17} - \frac{5}{17}i$$

**Example 37.** Division komplexer Zahlen

**Example.**

$$\begin{aligned} \frac{3 + 4i}{1 + 2i} &= \frac{(3 + 4i)(1 - 2i)}{(1 + 2i)(1 - 2i)} \\ &= \frac{3 - 6i + 4i - 8}{1^2 + 2^2} \\ &= \frac{11 - 2i}{5} \end{aligned}$$

**Example 38.** Was ist  $z^{-1}$ ,  $z \in \mathbb{C}$ , also was ist das Ergebnis der Gleichung  $(a + bi) \cdot x = 1$ ?

**Example.**

$$\begin{aligned}(a + bi) \cdot x &= 1 \\ (a - bi)(a + bi) \cdot x &= (a - bi) \\ (a^2 + b^2) \cdot x &= a - bi \\ x &= \frac{a}{a^2 + b^2} - \frac{b}{a^2 + b^2}i = z^{-1}\end{aligned}$$

oder symbolisch ausgedrückt unter Verwendung von  $z \cdot \bar{z} = |z|^2$ :

$$\begin{aligned}zx &= 1 \\ z \cdot \bar{z} \cdot x &= \bar{z} \\ |z|^2 x &= \bar{z} \\ x &= \frac{\bar{z}}{|z|^2} = z^{-1}\end{aligned}$$

Damit besitzt jede komplexe Zahl  $\neq 0$  ein inverses Element  $\frac{1}{z}$ .

**14.4. Weitere Regeln für komplexe Zahlen.** Sei  $z = a + bi$ . Dann ist:

- $\overline{\bar{z}} = z$
- $\operatorname{Re}(z) = \frac{1}{2}(z + \bar{z})$
- $\operatorname{Im}(z) = \frac{1}{2i}(z - \bar{z})$
- Beweis:  $\frac{1}{2i}(z - \bar{z}) = \frac{1}{2i}(a + bi - (a - bi)) = \frac{1}{2i}2bi = b$
- $z\bar{z} = 0 \Leftrightarrow z = 0$ 
  - erster Teilbeweis:  $z\bar{z} = 0 \Rightarrow z = 0$

$$\begin{aligned}z\bar{z} = a^2 + b^2 &= 0 \\ \Rightarrow a = b &= 0\end{aligned}$$

- zweiter Teilbeweis:  $z\bar{z} = 0 \Leftarrow z = 0$  (trivial)

$$(15) \quad \overline{z_1 + z_2} = \bar{z}_1 + \bar{z}_2$$

Beweis durch Nachrechnen.

$$(16) \quad \overline{z_1 \cdot z_2} = \bar{z}_1 \cdot \bar{z}_2$$

Beweis durch Nachrechnen.

#### 14.5. Fundamentalsatz der Algebra.

**Definition 62.** Polynom

**Definition.** Eine Ausdruck  $a_n x^n + a_{n-1} x^{n-1} + \dots + a_0$  heißt:

- komplexes Polynom vom Grad  $n$ , wenn  $a_i \in \mathbb{C}$  und  $a_n \neq 0$ .
- reelles Polynom vom Grade  $n$ , wenn  $a_i \in \mathbb{R}$  und  $a_n \neq 0$ . Reelle Polynome sind eine Teilmenge der komplexen Polynome.

**Definition 63.** Polynomfunktion

**Definition.** Eine Funktion  $p(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_0$ , deren rechter Teil ein Polynom ist, heißt Polynomfunktion. Polynomfunktionen heißen auch ganzrationale Funktionen. Funktionen der Form  $\frac{a_n x^n + a_{n-1} x^{n-1} + \dots + a_0}{b_n x^n + b_{n-1} x^{n-1} + \dots + b_0}$  heißen auch gebrochenrationale Funktionen.

**Definition 64.** Nullstelle einer Polynomfunktion  $p(x)$

**Definition.**  $z \in \mathbb{C}$  heißt Nullstelle von  $p(x)$ , wenn  $p(z) = 0$ .

**Definition 65.** Fundamentalsatz der Algebra

**Definition.** Sei  $p(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_0$  ein komplexes Polynom vom Grad  $n$ . Es gilt:

- $p(x)$  besitzt mindestens eine Nullstelle  $z_1 \in \mathbb{C}$ .
- Ist  $z_1$  eine Nullstelle von  $p(x)$ , dann gilt  $p(x) = (x - z_1) \cdot q(x)$ , wobei  $q(x)$  ein komplexes Polynom vom Grad  $n - 1$  ist.  $(x - z_1)$  heißt Linearfaktor.

Konsequenz: Der Fundamentalsatz der Algebra kann fortgesetzt zur Abspaltung von Linearfaktoren  $(x - z_i)$  angewandt werden, da  $q(x)$  selbst ein komplexes Polynom ist. Die Abspaltung endet bei Polynomen vom Grad 0, das sind die komplexen Zahlen selbst. Also:

**Definition 66.** Konsequenz aus dem Fundamentalsatz der Algebra<sup>7</sup>

**Definition.**  $p(x)$  zerfällt in  $n$  Linearfaktoren. Genauer:

$$(17) \quad p(x) = a_n \cdot (x - z_1) \cdot (x - z_2) \cdot \dots \cdot (x - z_n)$$

wobei  $z_1, \dots, z_n$  Nullstellen von  $p(x)$  sind.

**Definition 67.**  $k$ -fache Nullstelle von  $p(x)$

**Definition.**  $z$  heißt  $k$ -fache Nullstelle von  $p(x)$ , wenn  $z$  unter den Zahlen  $z_1, \dots, z_n$  aus Gleichung 17 genau  $k$ -mal vorkommt.

**Definition 68.** Anzahl der Nullstellen eines komplexen Polynoms<sup>8</sup>

**Definition.** Ein komplexes Polynom vom Grad  $n$  besitzt genau  $n$  komplexe Nullstellen, wenn jede Nullstelle mit ihrer Vielfachheit gezählt wird. Sind  $z_1, z_2, \dots, z_r$  komplexe Nullstellen von  $p(x)$ , dann gilt:

$$p(x) = (x - z_1) \cdot (x - z_2) \cdot \dots \cdot (x - z_r) \cdot q(x)$$

wobei  $q(x)$  ein Polynom vom Grad  $n - r$  ist.

**Definition 69.** Komplexe Zahl und ihre konjugiert komplexe Zahl sind Nullstellen

**Definition.** Sei  $p(x)$  ein reelles Polynom. Ist  $z = a + bi$ ,  $b \neq 0$  eine komplexe Nullstelle von  $p(x)$ , so auch  $\bar{z} = a - bi$ . Beweis: Man setze alle Nullstellen ein:

$$a_n z^n + a_{n-1} z^{n-1} + \dots + a_0 = 0$$

Mit Gleichungen 15 und 16 folgt:

$$\begin{aligned} \Rightarrow \overline{a_n z^n + a_{n-1} z^{n-1} + \dots + a_0} &= \bar{0} \\ \Rightarrow a_n \bar{z}^n + a_{n-1} \bar{z}^{n-1} + \dots + a_0 &= 0 \end{aligned}$$

qed.

**Example 39.** Nullstellenbestimmung bei reellen Polynomen

**Example.** Nullstellen bei Polynomen vom Grad 3 und höher findet man durch experimentelles Einsetzen von Zahlen<sup>9</sup>. Es sei:

$$p(x) = 2x^4 - 16x^3 + 52x^2 - 80x + 48$$

(1) Durch Probieren findet man  $z_1 = 2$  als Nullstelle von  $p(x)$ . Polynomdivision:

$$(2x^4 - 16x^3 + 52x^2 - 80x + 48) : (x - 2) = 2x^3 - 12x^2 + 28x - 24 =: q_1(x)$$

(2) Durch Probieren findet man  $z_2 = 2$  ist eine Nullstelle von  $q_1(x)$ <sup>10</sup>. Polynomdivision:

$$(2x^3 - 12x^2 + 28x - 24) : (x - 2) = 2x^2 - 8x + 12 := q_2(x)$$

(3) Die Nullstellen des quadratischen reellen Polynoms  $q_2(x)$  können durch  $pq$ -Formel bestimmt werden:

$$\begin{aligned} 2x^2 - 8x + 12 &= 0 & | : 2 \\ \Rightarrow x^2 - 4x + 6 &= 0 \\ \Rightarrow z_{3|4} &= 2 \pm \sqrt{4 - 6} \\ &= 2 \pm \sqrt{2}i \\ \Rightarrow z_3 &= 2 + \sqrt{2}i \\ \vee z_4 &= 2 - \sqrt{2}i \end{aligned}$$

Damit ist  $q_2(x) = 2 \cdot (x - z_3) \cdot (x - z_4)$ . Beachte: Ein konstanter Faktor, durch den man das Polynom während der Nullstellenbestimmung dividiert, was sowohl bei  $pq$ - als auch bei  $abc$ -Formel geschieht, muss berücksichtigt werden! Es ist der höchste Koeffizient des Polynoms  $p(x)$ , hier 2.

<sup>7</sup>Auch oft selbst als »Fundamentalsatz der Algebra« bezeichnet.

<sup>8</sup>Auch oft selbst als »Fundamentalsatz der Algebra« bezeichnet.

<sup>9</sup>Vereinbarung: Diese Zahlen liegen in dieser Veranstaltung zwischen  $-3$  und  $3$ .

<sup>10</sup>Nullstellen von Polynomen dritten Grades können auch mit den Kardan'schen Formeln bestimmt werden. In dieser Veranstaltung werden sie ebenfalls erraten.

Also ist:

$$\begin{aligned} p(x) &= 2 \cdot (x - z_1) \cdot (x - z_2) \cdot (x - z_3) \cdot (x - z_4) \\ &= 2 \cdot (x - 2) \cdot (x - 2) \cdot (x - 2 - \sqrt{2}i) \cdot (x - 2 + \sqrt{2}i) \end{aligned}$$

Will man nur eine reelle Zerlegung, ergibt sich aus zwei komplexen Linearfaktoren jeweils ein quadratischer reeller Faktor:

$$p(x) = 2(x - 2)(x - 2)(x^2 - 4x + 6)$$

**Definition 70.** Zerlegung von Polynomen

**Definition.** Man kann immer zwei konjugiert komplexe Linearfaktoren zu einem reellen quadratischen Faktor zusammenfassen. Somit zerfällt jedes reelle Polynom in ein Produkt aus reellen quadratischen Polynomen und reellen Linearfaktoren.

**14.6. Darstellungsformen komplexer Zahlen in der Gaußschen Zahlenebene.** Man verwendet in der Praxis drei Formen der Darstellung komplexer Zahlen:

**kartesische Form:**  $a + bi$ , wie bisher verwendet.

**trigonometrische Form:**  $r \cdot (\cos \phi + i \sin \phi)$ , wobei  $\phi$  das Argument komplexen Zahl ist.

**exponentielle Form:** (auch: Euler-Form)  $re^{\phi}$ , wobei  $\phi$  das Argument komplexen Zahl ist.

**14.6.1. Darstellung in kartesischer Form.** Die Gaußsche Zahlenebene (auch: Gauß-Ebene) ist ein Koordinatensystem mit reeller und imaginärer Achse.  $z = a + bi$  wird als Zeiger (Vektor) von  $(0;0)$  zu  $(a;b)$  in der Gaußschen Zahlenebene dargestellt (vgl. Abbildung 23). Die Addition und Subtraktion komplexer Zahlen

ABBILDUNG 23. Darstellung von  $z = 2 + 3i$  in der Gauß-Ebene und Addition

entspricht der Addition und Subtraktion von Vektoren: man hängt die Vektoren durch Parallelverschiebung aneinander und erhält die Summe als Vektor von  $(0;0)$  zur Spitze des verschobenen Vektors  $z'_2$  (vgl. Abbildung 23). Die Reihenfolge der Verkettung ist irrelevant.

**Example 40.** Subtraktion komplexer Zahlen

**Example.** Die Subtraktion ist durch eine Addition darstellbar:  $z_1 - z_2 = z_1 + (-z_2)$ . Die Umkehrung des Vorzeichens ist eine Punktspiegelung in  $(0;0)$ , es folgt eine Addition.

ABBILDUNG 24. Subtraktion  $z_1 - z_2 = z_1 + (-z_2)$

**14.6.2. Darstellung in Polarkoordinaten und trigonometrischer Form.** Für Multiplikation, Division und Wurzelziehen eignet sich die kartesische Form nicht. Man stellt komplexe Zahlen deshalb in Polarkoordinaten und trigonometrischer Form dar, ebenfalls in der Gaußschen Zahlenebene.

ABBILDUNG 25. Polarform komplexer Zahlen

Bezeichnungen in Abbildung 25:

$\phi$ : Argument von  $z$

$r$ : Länge des Zeigers, Betrag von  $z$

$$r = \sqrt{a^2 + b^2} = |z|$$

**Pol:** Bezeichnung für den Zeiger einer komplexen Zahl in Polarkoordinaten.

$(r; \phi)$ : Polarkoordinaten von  $z$

**Definition 71.** trigonometrische Form

**Definition.** Es gilt:

$$a = r \cdot \cos \phi$$

$$b = r \cdot \sin \phi$$

Herleitung der trigonometrischen Form:

$$z = a + bi = r \cdot \cos \phi + ir \cdot \sin \phi = r(\cos \phi + i \sin \phi)$$

$z = r(\cos \phi + i \sin \phi)$  heißt die trigonometrische Form von  $z$ .

**Definition 72.** Besonderheiten der trigonometrischen Form

**Definition.**

- $\phi$  ist nicht eindeutig bestimmt, da  $\cos \phi$  und  $\sin \phi$  die Periode  $2\pi$  besitzen. Das heißt: mit  $\phi$  ist auch  $\phi + k \cdot 2\pi$  für jedes  $k \in \mathbb{Z}$  ein Argument von  $z$ . Wenn  $0 \leq \phi < 2\pi$  gilt, heißt  $\phi$  der Hauptwert von  $z$ .
- Darstellung der Zahl  $Z = 0$ :  $r = 0$ ,  $\phi$  ist unbestimmt, weil jeder beliebige Winkel verwendet werden kann.

14.6.3. *Umwandlung von kartesischer in trigonometrische Form.* Sei  $z = a + bi$ . Dann ist

$$r = |z| = \sqrt{a^2 + b^2}$$

$$\phi = \arctan \frac{b}{a}$$

Tangens ist mit Periode  $\pi$  periodisch, seine Definitionsmenge wird zur Bildung einer Umkehrfunktion also auf  $(-\frac{\pi}{2}; \frac{\pi}{2})$  eingeschränkt. Dieses Intervall ist damit auch die Bildmenge von  $\arctan(x)$ . Um Werte  $0 \leq \phi < 2\pi$  zu erhalten, muss also umgeformt werden:

- (1) Man berechne

$$\phi' = \arctan \left| \frac{b}{a} \right|$$

und erhält aufgrund der Betragsstriche einen positiven Winkel  $0 \leq \phi' < \frac{\pi}{2}$ .

- (2) Man ermittle mit einer Skizze, in welchem Quadranten  $z$  liegt und erhält  $\phi$  aus  $\phi'$  dann gemäß Abbildung 26.

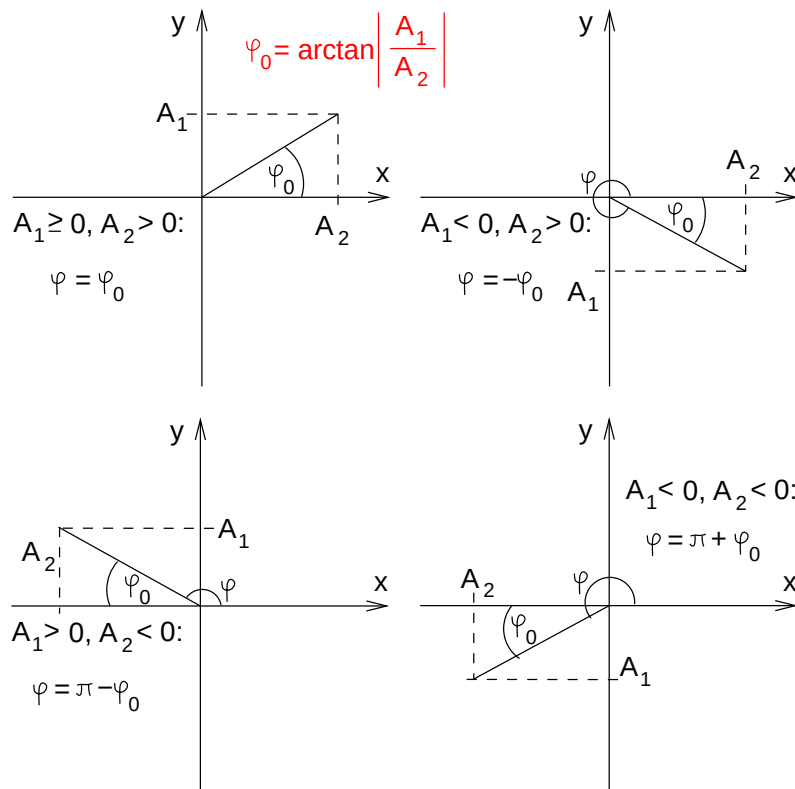


ABBILDUNG 26. Zu Umwandlung von kartesischer in trigonometrische Form

**Example 41.** Umwandlung von kartesischer in trigonometrische Form

**Example.** Sei  $z = \frac{3}{2} - i\frac{1}{2}\sqrt{3}$ . Vgl. Abbildung 26. Zuerst berechnet man

ABBILDUNG 27. Zur trigonometrischen Form von  $z = \frac{3}{2} - i\frac{1}{2}\sqrt{3}$

$$\phi' = \arctan \left| \frac{b}{a} \right| = \arctan \left| \frac{\frac{1}{2}\sqrt{3}}{\frac{3}{2}} \right| = \frac{\pi}{6}$$

Nun erhält man  $\phi = 2\pi - \phi' = \frac{11\pi}{6} = \arg(z)$ . Der Betrag ist:

$$r = \sqrt{a^2 + b^2} = \sqrt{\frac{9}{4} + \frac{3}{4}} = \sqrt{3}$$

Damit ist

$$z = r(\cos \phi + i \sin \phi) = \sqrt{3} \left( \cos \frac{\pi}{6} + i \sin \frac{\pi}{6} \right)$$

**Example 42.** Umwandlung von kartesischer in trigonometrische Form

**Example.**

•

$$\begin{aligned} z_0 &= 2 + 2i \\ r &= \sqrt{a^2 + b^2} = \sqrt{8} = 2\sqrt{2} \\ \phi &= \arctan \frac{b}{a} = \arctan 1 = \frac{\pi}{4} \\ z_0 &= 2\sqrt{2} \left( \cos \frac{\pi}{4} + i \sin \frac{\pi}{4} \right) \end{aligned}$$

•

$$\begin{aligned} z_1 &= -2 + 2i \\ r &= \sqrt{8} = 2\sqrt{2} \\ \phi' &= \arctan \left| \frac{b}{a} \right| = \frac{\pi}{4} \\ \phi &= \frac{3\pi}{4} \\ z_1 &= 2\sqrt{2} \left( \cos \frac{3\pi}{4} + i \sin \frac{3\pi}{4} \right) \end{aligned}$$

•

$$z_2 = -2 - 2i = 2\sqrt{2} \left( \cos \frac{5\pi}{4} + i \sin \frac{5\pi}{4} \right)$$

•

$$z_3 = 2 - 2i = 2\sqrt{2} \left( \cos \frac{7\pi}{4} + i \sin \frac{7\pi}{4} \right)$$

Umrechnung von Polarform in kartesische Form:

$$\begin{aligned} a &= r \cdot \cos \phi \\ b &= r \cdot \sin \phi \end{aligned}$$

#### 14.7. Multiplikation und Division komplexer Zahlen in trigonometrischer Form.

**Definition 73.** Multiplikation komplexer Zahlen in trigonometrischer Form

**Definition.** Seien  $z_1 = r_1(\cos \phi_1 + i \sin \phi_1)$ ,  $z_2 = r_2(\cos \phi_2 + i \sin \phi_2)$ . Dann ist das Produkt unter Verwendung der Additionstheoreme:

$$z_1 z_2 = r_1 r_2 (\cos \phi_1 \cos \phi_2 - \sin \phi_1 \sin \phi_2 + i(\cos \phi_1 \sin \phi_2 + \sin \phi_1 \cos \phi_2))$$

Diese Definition einer Multiplikation ist der wesentliche Unterschied zum  $\mathbb{R}^2$ .

$$\begin{aligned} (a + bi) \cdot (\alpha + \beta i) &= a\alpha + a\beta i + b\alpha i + b\beta i^2 \\ &= a\alpha + b\beta i^2 + (a\beta + \alpha b)i \end{aligned}$$

Unter Verwendung einer postulierten Definition  $i^2 = -1$  ergibt sich:

$$(a + bi) \cdot (\alpha + \beta i) = (a\alpha - b\beta) + (a\beta + \alpha b)i$$

Es kann durch ebensolchen Ausmultiplizieren gezeigt werden, dass für  $w, z \in \mathbb{C}$  gilt:  $z \cdot w = w \cdot z$ . Die Multiplikation ist also kommutativ. Weiter kann die Assoziativität der Multiplikation gezeigt werden, ebenso dass  $(1 + 0i)$  das neutrale Element der Multiplikation ist. Man kann also mit komplexen Zahlen als mit Paaren von reellen Zahlen rechnen.

Bei der Multiplikation komplexer Zahlen ordnet man dem Produkt  $z_1 z_2$  die Summe der Winkel  $\phi_1 + \phi_2$  zu. Diese Zuordnung ist von Logarithmen bekannt: der Winkel einer komplexen Zahl ist Teil ihres Logarithmus:

$$\begin{aligned} \ln z_1 &= \ln z + \arg(z) \cdot i \\ \ln -1 &= \ln 1 + \pi i \end{aligned}$$

**Example 43.** Rechenbeispiel für komplexe Zahlen

**Example.**

$$\begin{aligned}\overline{z + \bar{z}} &= \bar{z} + \bar{\bar{z}} \\ &= \bar{z} + z \\ &= z + \bar{z}\end{aligned}$$

**Example 44.** Rechenbeispiel für komplexe Zahlen

**Example.**

$$\begin{aligned}\overline{z \cdot \bar{z}} &= z \cdot \bar{z} \\ &= (a + bi) \cdot (a - bi) \\ &= a^2 + b^2\end{aligned}$$

Das Ergebnis ist also nach der dritten binomischen Formel und  $i^2 = -1$  reell. So kann gezeigt werden, dass das Ergebnis einer Rechnung reell ist.

**14.8. Betrag einer komplexen Zahl.** Per Definition ist ausgehend vom Satz des Pythagoras die Länge eines Pfeils zu einem Punkt der komplexen Zahlenebene der Betrag einer komplexen Zahl:

$$|z| = \sqrt{a^2 + b^2}$$

dabei ist bei einer beliebigen komplexen Zahl  $a^2 + b^2 \geq 0$ , so dass das Ziehen der Quadratwurzel mit den bisherigen Methoden ohne weiteres möglich ist und stets  $|z| \geq 0$ . Außerdem gilt:  $|z| = 0 \Leftrightarrow z = 0$ : ist die Länge einer komplexen Zahl 0, ist diese komplexe Zahl 0.

Es folgt mit 44:

$$(18) \quad |z|^2 = a^2 + b^2 = z \cdot \bar{z}$$

Betrag des Produktes  $z_1 \cdot z_2$ . Unter Verwendung von Gleichung 18 und Gleichung ?? gilt:

$$\begin{aligned}|z_1 \cdot z_2|^2 &= z_1 z_2 \cdot \overline{z_1 z_2} \\ &= z_1 \cdot z_2 \cdot \bar{z}_1 \cdot \bar{z}_2 \\ &= z_1 \cdot \bar{z}_1 \cdot z_2 \cdot \bar{z}_2 \\ &= |z_1|^2 \cdot |z_2|^2\end{aligned}$$

also folgt:

$$|z_1 \cdot z_2| = |z_1| \cdot |z_2|$$

Betrag der Summe  $z_1 + z_2$ .

$$\begin{aligned}|z_1 + z_2|^2 &= (z_1 + z_2) \cdot \overline{(z_1 + z_2)} \\ &= (z_1 + z_2) \cdot (\bar{z}_1 + \bar{z}_2) \\ &= z_1 \bar{z}_1 + z_1 \bar{z}_2 + z_2 \bar{z}_1 + z_2 \bar{z}_2 \\ &= |z_1|^2 + |z_2|^2 + (z_1 \bar{z}_2 + \bar{z}_1 z_2) \\ &\leq |z_1|^2 + |z_2|^2 + 2|z_1||z_2| = (|z_1| + |z_2|)^2\end{aligned}$$

Hier wurde verwendet, dass  $z_1 \bar{z}_2 + \bar{z}_1 z_2$  eine reelle Zahl sein muss. Dies ist die sogenannte Dreiecksungleichung, die formuliert werden kann als: die Summe von zwei Dreiecksseiten ist immer größer oder gleich der dritten Seite. Sie kann zur Abschätzung der Größe von Summen dienen.

Man kann also mit komplexen Zahlen rechnen wie mit reellen Zahlen; der einzige Unterschied ist, dass die komplexen Zahlen nicht nach ihrer Größe angeordnet werden können.

Für reelle Zahlen gilt die Eigenschaft NT (Nullteiler):  $a \cdot b = 0 \Rightarrow a = 0 \vee b = 0$ . Gilt dies auch für komplexe Zahlen? Ja, denn wenn  $z_1 \cdot z_2 = 0$ , so ist auch  $|z_1 \cdot z_2| = 0 \Leftrightarrow |z_1| \cdot |z_2| = 0 \Rightarrow |z_1| = 0 \vee |z_2| = 0$ .

Kann eine solche Eigenschaft NT auch für jedes Element aus einem beliebigen  $\mathbb{R}^n$  definiert werden? Für den  $\mathbb{R}^2$  konnte dies hier bei den komplexen Zahlen gezeigt werden.

Nun ist im  $\mathbb{R}^3$  eine ganze Reihe Multiplikationen definiert, aber weder die naive Multiplikation  $(a, b) \circ (c, d) = (a \cdot c, b \cdot d)$  noch das Kreuzprodukt haben den Nullteiler als geforderte Eigenschaft. Hier kann das Produkt nämlich auch 0 sein, ohne dass einer der Faktoren 0 ist. Ist es möglich, eine solche Möglichkeit im  $\mathbb{R}^3$  einzuführen? Nein, aber es funktioniert im  $\mathbb{R}^4$ , wo aber die Multiplikation aber nicht mehr kommutativ ist: Jede

Zahl kann zerlegt werden in vier Vektoren  $1, i, j, k$  (sog. Quaternionen), für die dann folgende Rechengesetze gelten

$$\begin{aligned} i &= (0, 1, 0, 0) \\ j &= (0, 0, 1, 0) \\ k &= (0, 0, 0, 1) \\ 1 &= (1, 0, 0, 0) \\ ij &= k \\ jk &= i \\ ki &= j \\ ij &= -ji \end{aligned}$$

Schließlich konnte durch Frank Adams 1963 bewiesen werden, dass nur im  $\mathbb{R}$ ,  $\mathbb{R}^2$ ,  $\mathbb{R}^4$ ,  $\mathbb{R}^8$  und  $\mathbb{R}^{16}$  eine solche Multiplikation eingeführt werden kann.

**14.9. Wurzeln aus komplexen Zahlen.** Es sei

$$\begin{aligned} z &= 3 + 4i = ((5; 0, 927295 \dots)) \\ w &= ((\sqrt{5}; 0, 463147 \dots)) \end{aligned}$$

So ist  $w^2 = ((\sqrt{5}^2; 2 \cdot 0, 463147)) = ((5; 0, 927295 \dots)) = z$ . So kann man nun durch Wurzelziehen aus dem Betrag und halbieren des Winkels sofort eine Quadratwurzel einer komplexen Zahl angeben, sofern diese in Polarkoordinaten vorliegt. Ebenso kann eine Wurzel  $n$ -ter Ordnung sofort angegeben werden: Ist  $z = ((r, \phi))$  gegeben und setzt man  $w = \sqrt[n]{z} = ((\sqrt[n]{r}; \frac{\phi}{n}))$ , so ist  $w^n = z$ .

**Example 45.** Lösung von  $x^2 = -1$

**Example.**

$$\begin{aligned} x^2 &= -1 = (1, \pi) \\ x &= ((1, \frac{\pi}{2})) = i \end{aligned}$$

**Example 46.** Lösung von  $x^2 = i$

**Example.**

$$\begin{aligned} x^2 &= i = ((1, \frac{\pi}{2})) \\ x &= ((1, \frac{\pi}{4})) \\ &= \frac{\sqrt{2}}{2} + i \cdot \frac{\sqrt{2}}{2} \end{aligned}$$

Die Darstellung  $\sin \frac{\pi}{4} = \cos \frac{\pi}{4} = \sqrt{\frac{1}{2}} = \frac{1 \cdot \sqrt{2}}{\sqrt{2} \cdot \sqrt{2}} = \frac{\sqrt{2}}{2}$  ergibt sich aus dem Satz des Pythagoras bei der Darstellung von  $\phi = \frac{\pi}{4}$  im Einheitskreis.

**14.10. komplexe Wurzeln aus 1.** Die Zahl 1 kann geschrieben werden als  $x_1^3 = ((1; 0))$ ,  $x_2^3 = ((1, 2\pi))$ ,  $x_3^3 = ((1, 4\pi))$ ,  $x_n^3 = ((1, n \cdot 2\pi))$ , d.h. es können Vollwinkel addiert werden, ohne die Zahl zu verändern. Damit ergeben sich nach obiger Darstellung jedoch unterschiedliche komplexe Wurzeln:

$$\begin{aligned} x_1 &= ((1; 0)) \\ x_2 &= ((1; \frac{2\pi}{3})) \\ x_3 &= ((1; \frac{4\pi}{3})) \\ x_4 &= ((1; \frac{6\pi}{3})) = ((1; 2\pi)) = x_1 \\ x_n &= ((1; \frac{n \cdot 2\pi}{3})) \end{aligned}$$

Ab  $x_4$  wiederholen sich die Zahlenwerte der Wurzeln, so dass  $x_1 \dots x_3$  die komplexen Wurzeln der Zahl 1 sind. Stellt man sie im Einheitskreis dar, so entsteht ein gleichzeitiges Dreieck, dessen eine Ecke auf  $x_1 = ((1; 0))$  liegt.



Was ergibt sich nun für die Lösungen von  $x^4 = 1$ ?

$$\begin{aligned} x_1^4 &= ((1; 0)) \\ x_1 &= ((1; 0)) = 1 \\ x_2^4 &= ((1; 2\pi)) \\ x_2 &= ((1; \frac{\pi}{2})) = i \\ x_3^4 &= ((1; 4\pi)) \\ x_3 &= ((1; \pi)) = -1 \\ x_4^4 &= ((1; 6\pi)) \\ x_4 &= ((1; \frac{3}{2}\pi)) = -i \end{aligned}$$

Im Einheitskreis ergibt sich also für diese Wurzeln ein Quadrat mit einer Ecke in  $x_1$ .

Für die Lösungen von  $x^k = 1$  folgt daher: Sie liegen auf einem im Einheitskreis gezeichneten regelmäßigen  $k$ -Eck, dessen eine Ecke auf  $k_1 = 1$  liegt.

**Definition 74.**  $k$ -te Wurzel aus 1

**Definition.** Die Gleichung  $z^k = 1$  hat in den komplexen Zahlen die Lösungen  $((1; n\frac{2\pi}{k}))$  mit  $n = 0, 1, \dots, k-1$ , das sind die Potenzen von  $\zeta_k = ((1; \frac{2\pi}{k}))$ . Es gilt also:  $\zeta_k^k = 1$ ; alle Lösungen von  $z^k = 1$  lassen sich schreiben als  $\zeta_k^n$  mit  $n = 0, 1, \dots, k-1$ .

## 15. MATRIZEN

### 15.1. Definition einer Matrix.

**Definition 75.** Matrix

**Definition.** Eine rechteckige Anordnung von  $m \cdot n$  Elementen  $a_{ik}$  in  $m$  Zeilen und  $n$  Spalten heißt  $(m, n)$ -Matrix. Allgemein:

$$A = (a_{ik}) = \begin{pmatrix} a_{11} & a_{12} & a_{13} & \dots & a_{1n} \\ a_{21} & a_{22} & a_{23} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots & \dots \\ a_{m1} & a_{m2} & a_{m3} & \dots & a_{mn} \end{pmatrix}$$

### 15.2. Verknüpfungen von Matrizen.

**Definition 76.** Summe und Differenz von Matrizen

**Definition.** Seien  $A, B$   $(m, n)$ -Matrizen. Dann ist ihre Summe bzw. Differenz  $A \pm B = C$  definiert über komponentenweise Addition bzw. Subtraktion gemäß  $c_{ik} = a_{ik} \pm b_{ik}$ .

**Definition 77.** Skalarmultiplikation einer Matrix

**Definition.** Sei eine  $(m, n)$ -Matrix  $A$  und ein Skalar  $k$ . Dann ist die Skalarmultiplikation  $kA = C$  definiert als komponentenweise Multiplikation mit dem Skalar gemäß  $c_{ik} = k \cdot a_{ik}$ .

**Definition 78.** Produkt von Matrizen

**Definition.** Das Matrizenprodukt  $A \cdot B = C$  ist nur definiert, wenn die Spaltenzahl von  $A$  gleich der Zeilenzahl von  $B$  ist. Sei  $A$  eine  $(r, n)$ -Matrix und  $B$  eine  $(n, s)$ -Matrix, dann ist das Produkt  $A \cdot B = C$  eine  $(r, s)$ -Matrix gemäß

$$c_{ik} = a_{i1}b_{1k} + a_{i2}b_{2k} + \dots + a_{in}b_{nk}$$

Man multipliziert also jede Zeile von  $A$  mit jeder Spalte von  $B$  und schreibt das Ergebnis an die Stelle, die von der Zeilennummer von  $A$  und der Spaltennummer von  $B$  angegeben wird.

**Definition 79.** Falk-Schema der Matrizenmultiplikation

**Definition.** Das Falk-Schema ist eine übersichtliche Schreibweise zur Multiplikation  $A \cdot B = C$  von Matrizen. Man notiert  $A$  unten links,  $B$  oben rechts und erhält  $C$  unten rechts, indem man für jedes Element  $c_{ik}$  das Produkt der Zeile  $i$  von  $A$  mit der Spalte  $k$  von  $B$  bildet. Beispiel:

$$\begin{array}{c|cc} & 1 & 0 \\ & 0 & 1 \\ \hline 1 & 2 & 1 & 2 \\ 3 & 4 & 3 & 4 \\ 5 & 6 & 5 & 6 \end{array}$$

$$\Rightarrow \begin{pmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{pmatrix}$$

**Definition 80.** Assoziativgesetz der Matrixmultiplikation

**Definition.** Falls die Matrizenmultiplikation ausführbar ist, gilt:

$$(AB) \cdot C = A \cdot (BC)$$

Bemerkung: Die Matrizenmultiplikation ist jedoch nicht kommutativ!

**Definition 81.** Distributivgesetze der Matrizenrechnung

**Definition.** Falls die Matrizenmultiplikation ausführbar ist, gilt:

$$\begin{aligned} A \cdot (B + C) &= AB + AC \\ (B + C) \cdot D &= BD + CD \end{aligned}$$

**Definition 82.**  $(n, n)$ -Einheitsmatrix

**Definition.** Als  $(n, n)$ -Einheitsmatrix wird bezeichnet:

$$E_{n,n} := E_n := \begin{pmatrix} 1 & 0 & 0 & \dots & 0 \\ 0 & 1 & 0 & \dots & 0 \\ 0 & 0 & 1 & & 0 \\ \vdots & \vdots & & \ddots & \vdots \\ 0 & 0 & 0 & \dots & 1 \end{pmatrix}$$

$E_5$  bezeichnet also z.B. eine  $(5, 5)$ -Matrix, deren Hauptdiagonale mit 1 belegt ist und deren sonstige Elemente 0 sind.

**Definition 83.** Neutrales Element der Matrixmultiplikation

**Definition.** Für  $(m, n)$ -Matrizen  $A$  gilt:

$$A \cdot E_n = E_m \cdot A = A$$

**15.3. Reguläre und inverse Matrizen.** Für  $n$ -reihige quadratische Matrizen ( $n$  fest) ist das Produkt immer definiert. Es gelten die folgenden Axiome:

- V:** (Verknüpfungsaxiom) denn  $A \cdot B$  ist ebenfalls  $n$ -reihig
- A:** (Assoziativgesetz) denn es gilt  $(AB) \cdot C = A \cdot (BC)$
- N:** (Neutrales Element) denn es gilt  $E_n A = A E_n = A$  für jedes  $n$ -reihige  $A$
- D:** (Distributivgesetze)

**Definition 84.** reguläre Matrix, singuläre Matrix, inverse Matrix

**Definition.** Eine  $n$ -reihige Matrix heißt regulär, wenn es ein  $A'$  gibt mit  $A \cdot A' = E_n$ . Andernfalls heißt  $A$  singulär.  $A'$  heißt inverse Matrix zu  $A$ , in Zeichen  $A^{-1} := A'$ .

**Definition 85.** Gruppe der regulären  $n$ -reihigen Matrizen

**Definition.** Die Menge der regulären  $n$ -reihigen Matrizen bildet eine Gruppe bzgl. der Multiplikation. Das Kommutativgesetz (K) gilt im allgemeinen nicht, so dass sie keinen Körper bildet. Es gilt:

- V:** (Verknüpfungsaxiom) denn es gilt  $A, B$  sind regulär  $\Rightarrow A \cdot B$  ist regulär
- A:** (Assoziativgesetz) denn es gilt  $(AB) \cdot C = A \cdot (BC)$
- N:** (Neutrales Element) denn es gilt  $E_n A = A E_n = A$  für jedes  $n$ -reihige  $A$
- I:** (Inverses Element) denn es gilt  $A^{-1} A = E_n$

**Example 47.** reguläre Matrix

**Example.**

- Alle Einheitsmatrizen sind reguläre Matrizen, denn es gilt  $E_n \cdot E_n = E_n$ . Die Einheitsmatrix ist also identisch mit ihrer inversen Matrix:  $E_n = E_n^{-1}$ .
- Matrizen, deren Elemente außerhalb der Hauptdiagonalen 0 sind, auf der Hauptdiagonalen jedoch ungleich 0, sind ebenfalls reguläre Matrizen. Sei eine solche Matrix

$$D := \begin{pmatrix} a_{11} & 0 & \dots & 0 \\ 0 & a_{22} & & 0 \\ \vdots & & \ddots & \\ 0 & 0 & & a_{nn} \end{pmatrix}$$

dann ist ihre inverse Matrix

$$D^{-1} = \begin{pmatrix} \frac{1}{a_{11}} & 0 & \dots & 0 \\ 0 & \frac{1}{a_{22}} & & 0 \\ \vdots & & \ddots & \\ 0 & 0 & & \frac{1}{a_{nn}} \end{pmatrix}$$

denn es gilt

$$D \cdot D^{-1} = D^{-1} \cdot D = E$$

Es gibt eine Beziehung zwischen Matrizen und linearen Gleichungssystemen. Ein beliebiges lineares Gleichungssystem

$$\begin{aligned} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n &= b_1 \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n &= b_2 \\ &\vdots = \vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n &= b_m \end{aligned}$$

entspricht dem Produkt einer Koeffizientenmatrix  $A$  mit einem Unbekanntenvektor  $x$  zu einem Ergebnisvektor  $b$ :

$$\Leftrightarrow \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \dots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix} \cdot \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_m \end{pmatrix}$$

**Definition 86.** Lösbarkeit linearer  $(n, n)$ -Gleichungssysteme über reguläre Matrix

**Definition.** Sei eine Koeffizientenmatrix  $A = (a_{ik})_{n,n}$ . Sei der Ergebnisvektor  $b = \begin{pmatrix} b_1 \\ \vdots \\ b_n \end{pmatrix}$  ein beliebiger

reeller Spaltenvektor. Dann gilt (ohne Beweis):

$$A \text{ ist regulär} \Leftrightarrow Ax = b \text{ ist lösbar}$$

Lösbarkeit bedeutet hier: es existiert genau eine Lösung der Gleichung  $Ax = b$ , nämlich  $x = A^{-1} \cdot b$ . Nachweis durch Einsetzen:

$$\begin{aligned} Ax &= b, x = A^{-1}b \\ A(A^{-1}b) &= b \\ \Leftrightarrow (AA^{-1})b &= b \\ \Leftrightarrow Eb &= b \\ \Leftrightarrow b &= b \quad (w) \end{aligned}$$

**Definition 87.** Ermittlung der inversen Matrix

**Definition.** Sei eine Matrix  $A$ .  $A^{-1}$  bestehe aus den Spaltenvektoren  $A^{-1} = (x^{(1)}, x^{(2)}, \dots, x^{(n)})$ . Diese Spaltenvektoren erhält man aufgrund der Beziehung  $AA^{-1} = E$  durch Lösen der folgenden  $n$  linearen Gleichungssysteme:

$$\begin{aligned} Ax^{(1)} &= \begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix} \\ Ax^{(2)} &= \begin{pmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{pmatrix} \\ &\vdots \end{aligned}$$

$$Ax^{(n)} = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{pmatrix}$$

Diese Gleichungssysteme kann man im Gauß-Schema simultan lösen.

**Example 48.** Lösen eines linearen Gleichungssystems über inverse Matrix

**Example.** Was ist die Lösung  $x_1, x_2, x_3$  des folgenden linearen Gleichungssystems?

$$Ax = b \Leftrightarrow \begin{pmatrix} 0 & 1 & 3 \\ 1 & 2 & 0 \\ 1 & 1 & -1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}$$

Nach Definition 86 ist  $x = A^{-1}b$ . Ermittlung von  $A^{-1}$  im Gauß-Verfahren, hier abgekürzt:

|          |   |    |                |                |                |
|----------|---|----|----------------|----------------|----------------|
| 0        | 1 | 3  | 1              | 0              | 0              |
| 1        | 2 | 0  | 0              | 1              | 0              |
| 1        | 1 | -1 | 0              | 0              | 1              |
| $\vdots$ |   |    | $\vdots$       |                |                |
| 1        | 0 | 0  | 1              | -2             | 3              |
| 0        | 1 | 0  | $-\frac{1}{2}$ | $\frac{3}{2}$  | $-\frac{3}{2}$ |
| 0        | 0 | 1  | $\frac{1}{2}$  | $-\frac{1}{2}$ | $\frac{1}{2}$  |

$$\text{Mit dieser Lösung ist } x = A^{-1}b = \begin{pmatrix} 1 & -2 & 3 \\ -\frac{1}{2} & \frac{3}{2} & -\frac{3}{2} \\ \frac{1}{2} & -\frac{1}{2} & \frac{1}{2} \end{pmatrix} \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} = \begin{pmatrix} 6 \\ 7 \\ 1 \end{pmatrix}$$

#### 15.4. Transformation. <sup>11</sup>

**Definition 88.** Transformation

**Definition.** Die Bewegung eines Punktes  $P$ , dargestellt durch seinen Ortsvektor  $\vec{p}$ , kann durch eine Transformationsfunktion  $y(\vec{x})$  dargestellt werden, in der  $\vec{p}$  mit der Transformationsmatrix zu multiplizieren ist. Bezeichnungen:

- Die Transformation ist eine Skalierung

$$y(\vec{x}) = S \cdot \vec{x}$$

- Die Transformation ist eine Translation (Verschiebung)

$$y(\vec{x}) = T \cdot \vec{x}$$

- Die Transformation ist eine Drehung oder eine kombinierte Transformation

$$y(\vec{x}) = M \cdot \vec{x}$$

**Definition 89.** Drehung um eine Achse: anschauliche Definition

**Definition.** Eine positive Drehung eines Punktes  $P$  um einen Winkel  $\phi > 0$  um eine Achse  $g : \vec{X} = \vec{q} + \lambda \vec{a}$  wird wie folgt erklärt: Man blickt entgegen der Richtung von  $\vec{a}$  und dreht  $P$  gegen den Uhrzeigersinn in einer zu  $g$  senkrechten Ebene.

**Definition 90.** Inhomogene Koordinaten

**Definition.** Koordinaten der Form  $P(x, y, z)$  nennt man inhomogene Koordinatendarstellung von  $P$ .

**Definition 91.** Homogene Koordinaten

**Definition.** Koordinaten der Form  $P(xw, yw, zw, w)$  nennt man homogene Koordinatendarstellung von  $P$ .  $w \in \mathbb{R}$ , wir setzen  $w = 1$ , so dass Punkte wie in  $P(x, y, z, 1)$  dargestellt werden. Homogene Koordinaten sind notwendig, um Translationen ebenfalls durch eine Matrizenmultiplikation, nämlich mit der Translationsmatrix  $T$ , darstellen zu können.

##### 15.4.1. Translation in inhomogenen Koordinaten.

Aufgabe. Berechnen Sie den Punkt  $P'$  durch Verschiebung des Punktes  $P$  um einen Vektor  $\vec{t}$ .

Algorithmus.

$$\vec{p}' = \vec{p} + \vec{t}$$

<sup>11</sup>Hier folgt keine vollständige Darstellung des Stoffes, sondern nur eine Algorithmik zum Lösen von Aufgaben bestimmter Typen.

15.4.2. *Skalierung in inhomogenen Koordinaten.*

Aufgabe. Berechnen Sie die Skalierungsmatrix  $S$  zur Skalierung eines Objektes in inhomogenen Koordinaten um einen Vektor  $\vec{s}$ . Fixpunkt der Skalierung ist  $O$ .

Algorithmus.  $\vec{s}$  hat drei identische Komponenten  $s_x = s_y = s_z$ , die den Skalierungsfaktor angeben.

$$S = \begin{pmatrix} s_x & 0 & 0 \\ 0 & s_y & 0 \\ 0 & 0 & s_z \end{pmatrix}$$

15.4.3. *Drehung in inhomogenen Koordinaten.*

Aufgabe. Berechnen Sie die Drehmatrix  $M$  zur Drehung eines Objektes in inhomogenen Koordinaten um einen Winkel  $\phi$  um eine gegebene Achse. Ein Objekt wird gedreht, indem man alle Punkte des Objektes mit  $y(\vec{x}) = M \cdot \vec{x}$  dreht.

Algorithmus. Um welche Achse soll gedreht werden?

- Drehung um die  $x$ -Achse

$$M = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \phi & -\sin \phi \\ 0 & \sin \phi & \cos \phi \end{pmatrix} = X(\phi)$$

- Drehung um die  $y$ -Achse

$$M = \begin{pmatrix} \cos \phi & 0 & \sin \phi \\ 0 & 1 & 0 \\ -\sin \phi & 0 & \cos \phi \end{pmatrix} = Y(\phi)$$

- Drehung um die  $z$ -Achse

$$M = \begin{pmatrix} \cos \phi & -\sin \phi & 0 \\ \sin \phi & \cos \phi & 0 \\ 0 & 0 & 1 \end{pmatrix} = Z(\phi)$$

- Drehung um eine Gerade durch  $O$ :  $g: \vec{X} = \lambda \vec{e}$ ,  $|\vec{e}| = 1$

$$(19) \quad M = \vec{e} \cdot \vec{e}^T + (E_3 - \vec{e} \cdot \vec{e}^T) \cdot (\cos \phi + U \sin \phi) = A(\phi)$$

mit

$$\vec{e} \cdot \vec{e}^T = \begin{pmatrix} e_x \\ e_y \\ e_z \end{pmatrix} \cdot \begin{pmatrix} e_x & e_y & e_z \end{pmatrix} = \begin{pmatrix} e_x e_x & e_x e_y & e_x e_z \\ e_y e_x & e_y e_y & e_y e_z \\ e_z e_x & e_z e_y & e_z e_z \end{pmatrix}$$

$$E_3 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

$$U = \begin{pmatrix} 0 & -e_z & e_y \\ e_z & 0 & -e_x \\ -e_y & e_x & 0 \end{pmatrix}$$

- Drehung um eine beliebige Gerade  $g: \vec{X} = \vec{p}_0 + \lambda \vec{e}$ ;  $|\vec{e}| = 1$ ;  $P_0(x_0; y_0; z_0)$ . Diese Transformation kann in homogenen Koordinaten nicht allein durch eine Transformationsmatrix dargestellt werden, weil eine Verschiebung notwendig ist. Unter Verwendung von Gleichung 19 ergibt sich ein verschobener Punkt  $\vec{p}'$  aus  $\vec{p}$  durch

$$\vec{p}' = A(\phi) \cdot (\vec{p} - \vec{p}_0) + \vec{p}_0$$

15.4.4. *Translation in homogenen Koordinaten.*

Aufgabe. Berechnen Sie die Translationsmatrix  $T$  zur Translation eines Punktes in homogenen Koordinaten um einen Vektor  $\vec{t}$ .

Algorithmus.

$$(20) \quad T = T(t_x, t_y, t_z) = \begin{pmatrix} 1 & 0 & 0 & t_x \\ 0 & 1 & 0 & t_y \\ 0 & 0 & 1 & t_z \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

15.4.5. *Skalierung in homogenen Koordinaten.*

Aufgabe. Berechnen Sie die Skalierungsmatrix  $S$  zur Skalierung eines Objektes in inhomogenen Koordinaten um einen Vektor  $\vec{s}$ . Fixpunkt der Skalierung ist  $O$ .

Algorithmus.  $\vec{s}$  hat drei identische Komponenten  $s_x = s_y = s_z$ , die den Skalierungsfaktor angeben.

$$S = \begin{pmatrix} s_x & 0 & 0 & 0 \\ 0 & s_y & 0 & 0 \\ 0 & 0 & s_z & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

#### 15.4.6. Drehung in homogenen Koordinaten.

Aufgabe. Berechnen Sie die Drehmatrix  $M$  zur Drehung eines Objektes in homogenen Koordinaten um einen Winkel  $\phi$  um eine beliebige Gerade  $g: \vec{X} = \vec{p}_0 + \lambda \vec{e}$ ;  $|\vec{e}| = 1$ ;  $P_0(x_0; y_0; z_0; 1)$ .

Algorithmus. Sei  $A^*(\phi)$  die Drehmatrix bei Drehung in inhomogenen Koordinaten nach Gleichung 19. Dann ist die Drehmatrix  $M$ :

$$M = T(x_0, y_0, z_0) \cdot A(\phi) \cdot T(-x_0, -y_0, -z_0) = M(\phi, g)$$

mit

$$A(\phi) = \begin{pmatrix} A^*(\phi) & 0 \\ 0 & 1 \end{pmatrix}$$

$T(x_0, y_0, z_0)$ ,  $T(-x_0, -y_0, -z_0)$ : Translationsmatrizen in homogenen Koordinaten nach Gleichung 20.

#### 15.4.7. Kombinierte Transformation in homogenen Koordinaten.

Aufgabe. Mit einem Objekt sind verschiedene Translationen in bestimmter Reihenfolge durchzuführen. Berechnen Sie die Bewegungsmatrix  $M$  in homogenen Koordinaten für die Gesamtbewegung.

Algorithmus. Man Berechne die Matrizen der einzelnen Translationen und multipliziere sie miteinander, beginnend mit der zuletzt durchzuführenden Bewegung  $M_3$  bis zur zuerst durchzuführenden Bewegung  $M_1$ :

$$M = M_3 \cdot M_2 \cdot M_1$$

## 16. DETERMINANTEN

16.1. **Motivation.** Die allgemeine Lösung eines linearen  $(2, 2)$ -Gleichungssystems lässt sich errechnen zu

$$x_1 = \frac{b_1 a_{22} - b_2 a_{12}}{a_{11} a_{22} - a_{21} a_{12}}$$

$$x_2 = \frac{b_2 a_{11} - b_1 a_{21}}{a_{22} a_{11} - a_{12} a_{21}}$$

Dabei ergeben sich die Nenner ja nach einem Kreuzschema aus der Koeffizientenmatrix:

$$A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \Rightarrow a_{11} a_{22} - a_{21} a_{12}$$

Welche Bedeutung haben dieses Schema und diese Zahlen, die Determinanten?

#### 16.2. Ermittlung von Determinanten.

**Definition 92.** Determinante einer  $(2, 2)$ -Matrix

**Definition.** Die Determinante der Matrix  $A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$  ist die Zahl  $a_{11} a_{22} - a_{12} a_{21}$ . In Zeichen:

$$\det A := |A| = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix}$$

$$:= a_{11} a_{22} - a_{12} a_{21}$$

**Definition 93.** Cramersche Regel

**Definition.** Die Lösung eines linearen  $(2, 2)$ -Gleichungssystems ist

$$x_1 = \frac{\begin{vmatrix} b_1 & a_{12} \\ b_2 & a_{22} \end{vmatrix}}{\begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix}}$$

$$x_2 = \frac{\begin{vmatrix} a_{11} & b_1 \\ a_{21} & b_2 \end{vmatrix}}{\begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix}}$$

Wenn die Nennerdeterminante  $\neq 0$  ist, so hat das Gleichungssystem genau eine Lösung.

**Definition 94.** Determinante

**Definition.** Dies ist eine Verallgemeinerung von Definition 92. Sei

$$A = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & & & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix}, m > 2$$

Sei  $U_{ij}$  jene Matrix, die aus  $A$  durch Streichen der  $i$ -ten Zeile und  $j$ -ten Spalte entsteht. Als Determinante von  $A$  wird definiert:

$$\begin{aligned} \det A &:= |A| \\ &:= \sum_{j=1}^n (-1)^{1+j} \cdot a_{1j} \cdot |U_{1j}| \\ &= a_{11} |U_{11}| - a_{12} |U_{12}| + a_{13} |U_{13}| - \dots \pm a_{1n} |U_{1n}| \end{aligned}$$

**Example 49.** allgemeine Bestimmung der Determinante einer  $(3, 3)$ -Matrix

**Example.** Es ist  $n = 3$ .

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} = a_{11} \begin{vmatrix} a_{22} & a_{23} \\ a_{32} & a_{33} \end{vmatrix} - a_{12} \begin{vmatrix} a_{21} & a_{23} \\ a_{31} & a_{33} \end{vmatrix} + a_{13} \begin{vmatrix} a_{21} & a_{22} \\ a_{31} & a_{32} \end{vmatrix}$$

**Example 50.** konkrete Bestimmung der Determinante einer  $(3, 3)$ -Matrix

**Example.**

$$\begin{aligned} \begin{vmatrix} 1 & 2 & 3 \\ 0 & 3 & 0 \\ 1 & 1 & -4 \end{vmatrix} &= 1 \begin{vmatrix} 3 & 0 \\ 1 & -4 \end{vmatrix} - 2 \begin{vmatrix} 0 & 0 \\ 1 & -4 \end{vmatrix} + 3 \begin{vmatrix} 0 & 3 \\ 1 & 1 \end{vmatrix} \\ &= 1 \cdot (-12) - 2 \cdot (0) + 3 \cdot (-3) \\ &= -21 \end{aligned}$$

**Definition 95.** Entwicklungssatz von Laplace (ohne Beweis)

**Definition.** In einer Determinante  $|A|$  kann man die erste Zeile mit einer beliebigen anderen tauschen und also zur  $i$ -ten statt nur zur 1-ten Zeile entwickeln. Der geringste Rechenaufwand ergibt sich, wenn man nach der Zeile oder Spalte mit den meisten Nullen entwickelt.

- Entwickeln nach der  $i$ -ten Zeile ( $i = 1, 2, 3, \dots, n$ ):

$$|A| = \sum_{j=1}^n (-1)^{i+j} a_{ij} \cdot |U_{ij}|$$

- Entwickeln nach der  $j$ -ten Spalten ( $j = 1, 2, 3, \dots, n$ ):

$$|A| = \sum_{i=1}^n (-1)^{i+j} a_{ij} \cdot |U_{ij}|$$

**Example 51.** Determinante bestimmen nach dem Entwicklungssatz von Laplace

**Example.** Wir entwickeln stets zu der Zeile oder Spalte mit den meisten Nullen, hier also zuerst zur zweiten Zeile:

$$\begin{aligned} \begin{vmatrix} 1 & 2 & 0 & 4 \\ 0 & 1 & 2 & 0 \\ 3 & 0 & 1 & 0 \\ 2 & 1 & 0 & 1 \end{vmatrix} &= -0 \cdot \begin{vmatrix} 2 & 0 & 4 \\ 0 & 1 & 0 \\ 1 & 0 & 1 \end{vmatrix} + 1 \cdot \begin{vmatrix} 1 & 0 & 4 \\ 3 & 1 & 0 \\ 2 & 0 & 1 \end{vmatrix} - 2 \cdot \begin{vmatrix} 1 & 2 & 4 \\ 3 & 0 & 0 \\ 2 & 1 & 1 \end{vmatrix} + 0 \cdot \begin{vmatrix} 1 & 2 & 0 \\ 3 & 0 & 1 \\ 2 & 1 & 0 \end{vmatrix} \\ &= 1 \cdot \begin{vmatrix} 1 & 4 \\ 2 & 1 \end{vmatrix} - 2 \cdot \left( (-3) \cdot \begin{vmatrix} 2 & 4 \\ 1 & 1 \end{vmatrix} \right) \\ &= (1 - 8) - 2 \cdot (-3) \cdot (2 - 4) \\ &= -19 \end{aligned}$$

**Definition 96.** Determinante einer oberen Dreiecksmatrix

**Definition.** Sei  $A$  eine obere Dreiecksmatrix

$$A = \begin{pmatrix} a_{11} & \dots & \dots & a_{1n} \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & a_{nn} \end{pmatrix}$$

Dann ist

$$|A| = a_{11} \cdot a_{22} \cdot \dots \cdot a_{nn} = \prod_{i=1}^n a_{ii}$$

Beweis zu Definition 96. Sukzessive Entwicklung, stets nach der ersten Spalte, liefert:

$$\begin{aligned} |A| &= a_{11} \begin{vmatrix} a_{22} & \dots & \dots & a_{2n} \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & a_{nn} \end{vmatrix} = a_{11} a_{22} \begin{vmatrix} a_{33} & \dots & \dots & a_{3n} \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & a_{nn} \end{vmatrix} \\ &= \dots = a_{11} \cdot a_{22} \cdot \dots \cdot a_{nn} \end{aligned}$$

qed.

**Definition 97.** Regel von Sarrus

**Definition.** Die Regel von Sarrus ist ein Schema zur Bestimmung nur von dreieihigen Determinanten. Man wiederholt die ersten beiden Spalten rechts neben  $A$  und zeichnet alle Diagonalen mit drei Elementen ein. Man berechnet alle Produkte aus den Elementen jeweils einer Diagonale. Man addiert die Produkte, die aus den nach rechts unten laufenden Diagonalen gebildet wurden, und subtrahiert die Produkte, die aus den nach links unten laufenden Diagonalen gebildet wurden. Man erhält:

$$|A| = a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} - a_{31}a_{22}a_{13} - a_{32}a_{23}a_{11} - a_{33}a_{21}a_{12}$$

**Definition 98.** Bedeutung der Determinante für ein lineares  $(n, n)$ -Gleichungssystem

**Definition.** Sei  $Ax = b$  ein lineares Gleichungssystem mit  $n$  Zeilen und  $n$  Unbekannten. Es gilt:

$$\begin{aligned} |A| \neq 0 &\Leftrightarrow A \text{ ist regulär} \\ &\Leftrightarrow A \text{ ist invertierbar} \\ &\Leftrightarrow Ax = b \text{ ist eindeutig lösbar} \end{aligned}$$

**Definition 99.** Cramersche Regel für ein lineares  $(n, n)$ -Gleichungssystem

**Definition.** Sei  $Ax = b$  ein lineares Gleichungssystem mit  $n$  Zeilen und  $n$  Unbekannten.  $A$  sei eine invertierbare Matrix  $A = (a_1, \dots, a_n) \in \mathbb{R}^{n \times n}$ . Es gilt (ohne Beweis):

$$x_i = \frac{\det(a_1, \dots, a_{i-1}, b, a_{i+1}, \dots, a_n)}{\det |A|}$$

(Im Nenner wird die  $i$ -te Spalte von  $A$  durch  $b$  ersetzt,  $1 \leq i \leq n$ )

### 16.3. Rechenregeln für Determinanten.

**Definition 100.** Wirkung sogenannter »elementarer« Zeilen- und Spaltenoperationen auf Determinanten

**Definition.** Sei  $A = (a_{ik})_{m,n}$ . Wenn  $\tilde{A}$  aus  $A$  entsteht durch ...

- ... Vertauschung zweier Zeilen [Spalten], dann entsteht die Determinante durch  $\det \tilde{A} = -\det A$ .
- ... Addition des  $r$ -fachen einer Zeile [Spalte] zu einer anderen Zeile [Spalte], dann ändert sich die Determinante nicht:  $\det \tilde{A} = \det A$ .
- ... Multiplikation einer Zeile [Spalte] mit einem  $\lambda \in \mathbb{R} \setminus 0$ , dann ändert sich die Determinante entsprechend  $\det \tilde{A} = \lambda \cdot \det A$ .

(ohne Beweis)

Zur Berechnung von  $\det A$  bei  $n > 3$  ist es sinnvoll, zuerst mit den elementaren Zeilen- und Spaltenoperationen gemäß Definition 100 möglichst viele Nullen in einer Zeile oder Spalte zu erzeugen.

**Definition 101.** Multiplikationssatz für Determinanten

**Definition.** Seien  $A, B$  Matrizen vom Typ  $(n, n)$ . Es gilt:  $|A \cdot B| = |A| \cdot |B|$ .

**Example 52.** Multiplikationssatz für Determinanten



**Example.**

$$\left| \begin{pmatrix} 1 & 0 & 0 \\ 0 & 2 & 1 \\ 0 & 0 & 3 \end{pmatrix} \cdot \begin{pmatrix} 3 & 2 & 0 \\ 0 & 1 & 0 \\ 2 & 3 & 5 \end{pmatrix} \right| = 1 \cdot 6 \cdot 1 \cdot \left| \begin{pmatrix} 3 & 0 \\ 2 & 5 \end{pmatrix} \right| = 6 \cdot (15 - 0) = 90$$

**Definition 102.** Symmetrie in Zeilen und Spalten bei Determinanten

**Definition.** Seien  $A, B$  Matrizen vom Typ  $(n, n)$ . Es gilt:  $|A^T| = |A|$ .

## 17. VEKTOREN

**17.1. Vektoren.** Ein Vektor ist ein Pfeil, er hat immer eine bestimmte Richtung und eine bestimmte Länge. Jeder Vektor hat Anfang und Spitze.

**17.1.1. Koordinaten und Komponenten.**

**Koordinaten:** Die einzelnen Zahlen innerhalb des Vektors. Die Koordinaten des Vektors  $\begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}$  sind 1, 2 und 3.

**Komponenten:** Die einzelnen Teilvektoren eines Vektors. Die Komponenten des Vektors  $\begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}$  sind  $\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$ ,  $\begin{pmatrix} 0 \\ 2 \\ 0 \end{pmatrix}$  und  $\begin{pmatrix} 0 \\ 0 \\ 3 \end{pmatrix}$ .

**17.1.2. Länge eines Vektors.** Die Länge eines Vektors wird auch als Betrag dieses Vektors bezeichnet. Man berechnet den Betrag folgendermaßen:  $|\vec{v}| = \left| \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} \right| = \sqrt{v_1^2 + v_2^2 + v_3^2}$ .

**Example 53.** Berechne die Länge des Vektors  $\begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}$

**Example.** Die Länge des Vektors  $\begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}$  ist  $\sqrt{1^2 + 2^2 + 3^2} = \sqrt{14}$ .

**17.1.3. Summe zweier Vektoren.** Die Koordinaten des entstehenden Vektors entstehen durch Addition der Koordinaten der beiden Summanden.

**Example 54.**  $\begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + \begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix} = \begin{pmatrix} 4 \\ 4 \\ 4 \end{pmatrix}$

**17.1.4. Differenz zweier Vektoren.** Die Differenz zweier Vektoren berechnet sich in der selben Art und Weise wie die Addition. Der Unterschied besteht darin, dass die entstehenden Koordinaten durch Subtraktion der Koordinaten der ersten beiden Vektoren entstehen.

**Example 55.**  $\begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} - \begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix} = \begin{pmatrix} -2 \\ 0 \\ 2 \end{pmatrix}$

Geometrisch werden zwei Vektoren subtrahiert, indem man ihre Ausgangspunkte aneinanderlegt. Der Subtraktionsvektor reicht dann von der Spitze des zweiten zur Spitze des ersten Vektors. Zur Veranschaulichung dient [28](#)

**17.1.5. Skalarmultiplikation.** Die Multiplikation eines Vektors mit einer Zahl (einem »Skalar«) heißt Skalarmultiplikation. Jede Koordinate des Vektors wird dazu mit dem Skalar multipliziert.

**Example 56.**  $5 \cdot \begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix} = \begin{pmatrix} 15 \\ 10 \\ 5 \end{pmatrix}$

17.1.6. *Linearkombination.***Definition 103.** Linearkombination

**Definition.** Für  $n$  Vektoren  $\vec{a}_1, \vec{a}_2, \vec{a}_3, \dots, \vec{a}_n$  und  $n$  reelle Zahlen  $r_1, r_2, \dots, r_n$  heißt der Vektor  $\vec{b} = r_1 \cdot \vec{a}_1 + r_2 \cdot \vec{a}_2 + \dots + r_n \cdot \vec{a}_n$  eine Linearkombination der Vektoren  $\vec{a}_1, \vec{a}_2, \vec{a}_3, \dots, \vec{a}_n$ .

**Definition 104.** Lineare Abhängigkeit

**Definition.** Die Vektoren  $\vec{a}_1, \vec{a}_2, \vec{a}_3, \dots, \vec{a}_n$  heißen genau dann »linear abhängig«, wenn es Zahlen  $r_1, r_2, \dots, r_n$  aus dem reellen Bereich gibt, die nicht alle gleich 0 sind und für die  $r_1 \cdot \vec{a}_1 + r_2 \cdot \vec{a}_2 + \dots + r_n \cdot \vec{a}_n = 0$  gilt.

**Definition 105.** Lineare Unabhängigkeit

**Definition.** Die Vektoren  $\vec{a}_1, \vec{a}_2, \vec{a}_3, \dots, \vec{a}_n$  heißen genau dann »linear unabhängig«, wenn die Gleichung  $r_1 \cdot \vec{a}_1 + r_2 \cdot \vec{a}_2 + \dots + r_n \cdot \vec{a}_n = 0$  nur für  $r_1 = r_2 = \dots = r_n = 0$  gilt. Für andere Zahlen  $r_1, r_2, \dots, r_n$  muss  $r_1 \cdot \vec{a}_1 + r_2 \cdot \vec{a}_2 + \dots + r_n \cdot \vec{a}_n \neq 0$  gelten.

**Definition 106.** Lineare Abhängigkeit paralleler Vektoren

**Definition.** Zwei parallele Vektoren sind kollinear und immer linear abhängig, weil sie sich mit geeigneten Faktoren auslöschen, also 0 ergeben.

**Definition 107.** Komplanare Vektoren

**Definition.** Wenn 3 räumliche Vektoren linear abhängig sind, dann sind sie komplanar, d.h. sie liegen in einer Ebene.

17.1.7. *Basis.* Die Menge der Vektoren  $\{\vec{a}_1, \vec{a}_2, \vec{a}_3, \dots, \vec{a}_n\}$  von Vektoren des Vektorraumes  $V$  heißt Basis, wenn

- (1)  $\vec{a}_1, \vec{a}_2, \vec{a}_3, \dots, \vec{a}_n$  linear unabhängig sind
- (2) Sich jeder Vektor aus  $V$  als Linearkombination der Basisvektoren darstellen lässt:  $\vec{b} = r_1 \cdot \vec{a}_1 + r_2 \cdot \vec{a}_2 + r_3 \cdot \vec{a}_3 + \dots + r_n \cdot \vec{a}_n$

Bei einer *orthonormierten Basis* sind die Basisvektoren auf 1 normiert und stehen orthogonal zueinander, d.h. senkrecht.

17.1.8. *Dimension.* Die Anzahl der Vektoren, die in der Basis enthalten sind (also die Anzahl der Elemente in  $\{\vec{a}_1, \vec{a}_2, \vec{a}_3, \dots, \vec{a}_n\}$ ) heißt *Dimension*.

17.1.9. *Einheitsvektor.* Der Einheitsvektor  $\vec{e}_a$  eines Vektors  $\vec{a}$  hat die selbe Ausrichtung wie  $\vec{a}$ , die Länge des Einheitsvektors ist jedoch immer 1. Er berechnet sich mit  $\vec{e}_a = \vec{a} \cdot \frac{1}{|\vec{a}|} = \vec{a} \cdot \frac{1}{\sqrt{a_1^2 + a_2^2 + a_3^2}}$ .

**Example 57.** Einheitsvektor von  $\vec{v} = \begin{pmatrix} 2 \\ 5 \\ 3 \end{pmatrix}$

**Example.** Der Einheitsvektor ist  $\vec{e}_v = \begin{pmatrix} 2 \\ 5 \\ 3 \end{pmatrix} \cdot \frac{1}{\sqrt{38}}$ .

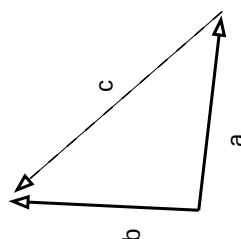


ABBILDUNG 28. Differenz zweier Vektoren

17.1.10. *Normalenvektor.* Der Normalenvektor eines Vektors steht senkrecht auf diesen. Seine Berechnung ist unterschiedlich für die zweite und dritte Dimension.

**Berechnung in  $\mathbb{R}^2$ :** Hier werden einfach die beiden Koordinaten vertauscht und bei einer das Vorzeichen umgedreht.

**Example 58.** Normalenvektor in  $\mathbb{R}^2$  von  $\vec{v} = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$

**Example.** Der Normalenvektor des Vektors  $\vec{v} = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$  ist  $\vec{v} = \begin{pmatrix} -2 \\ 1 \end{pmatrix}$ .

**Berechnung in  $\mathbb{R}^3$ :** Die Berechnung des Normalenvektors in  $\mathbb{R}^3$  ist etwas aufwendiger. Zum einen werden hier zwei Vektoren benötigt, um den Normalenvektor zu berechnen, da es im  $\mathbb{R}^3$  keinen einzelnen senkrecht stehenden Vektor zu einem anderen Vektor gibt. Dies ist für einen einzelnen Vektor immer eine gesamte Ebene. Zum anderen erfolgt die Berechnung mittels einer Matrix mit Zuhilfenahme der drei Einheitsvektoren. Der Vektor, der zu den Vektoren  $\vec{a} = \begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix}$  und  $\vec{b} = \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix}$  senkrecht steht ist

$$\vec{n} = \begin{vmatrix} e_1 & a_1 & b_1 \\ e_2 & a_2 & b_2 \\ e_3 & a_3 & b_3 \end{vmatrix} = e_1 \cdot (a_2 b_3 - a_3 b_2) + e_2 \cdot (a_3 b_1 - a_1 b_3) + e_3 \cdot (a_1 b_2 - a_2 b_1)$$

Es ist erlaubt aus dem berechneten Normalenvektor zu kürzen um diesen zu vereinfachen. Vereinfacht gesagt ist der Normalenvektor das Ergebnis des Spatproduktes  $\vec{n} = \vec{e} \cdot (\vec{a} \times \vec{b})$  der Vektoren  $\vec{a}$ ,  $\vec{b}$  und des Einheitsvektors  $\vec{e}$ . Siehe auch Kapitel 20.

**Example 59.** Normalenvektoren von  $\vec{a} = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}$  und  $\vec{b} = \begin{pmatrix} 4 \\ 1 \\ 2 \end{pmatrix}$

**Example.** Der Normalenvektor ist  $\vec{n} = \begin{pmatrix} 1 \\ 10 \\ -7 \end{pmatrix}$ .

17.2. **Punkte.** Der Vektor zwischen zwei Punkten berechnet sich als Differenz des zweiten Punkt und des ersten Punktes. Der Vektor zwischen  $A$  mit  $\vec{a} = \begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix}$  und  $B$  mit  $\vec{b} = \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix}$  ist also  $\vec{v} = \begin{pmatrix} b_1 - a_1 \\ b_2 - a_2 \\ b_3 - a_3 \end{pmatrix}$ .

**Example 60.** Der Vektor zwischen den Punkten  $A$  mit  $\vec{A} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$  und  $B$  mit  $\vec{B} = \begin{pmatrix} 2 \\ 3 \\ 1 \end{pmatrix}$  ist  $\vec{v} = \begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix}$ .

17.2.1. *Abstand zweier Punkte.* Der Abstand zweier Punkte berechnet sich als Betrag des Vektors, der zwischen diesen beiden Punkten liegt. Der Ausgangspunkt ist egal, da die Länge eines Vektors nicht von der Ausrichtung abhängt.

**Example 61.** Abstand zweier Punkte

**Example.** Die Punkte  $A \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$  und  $B \begin{pmatrix} 2 \\ 3 \\ 1 \end{pmatrix}$  haben einen Abstand von  $\sqrt{(1-2)^2 + (1-3)^2 + (1-1)^2} = \sqrt{5}$ .

## 18. SKALARPRODUKT

Das Skalarprodukt ist eine der beiden grundlegenden Arten zwei Vektoren zu multiplizieren. Wie der Name »Skalarprodukt« bereits sagt, ist das Ergebnis der Multiplikation eine Zahl.

**Koordinatendarstellung:** Das Skalarprodukt der Vektoren  $\vec{a} = \begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix}$  und  $\vec{b} = \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix}$  ist die Zahl

$$c = \vec{a} \cdot \vec{b} = a_1 \cdot b_1 + a_2 \cdot b_2 + a_3 \cdot b_3$$

**Vektordarstellung:** Das Skalarprodukt der Vektoren  $\vec{a} = \begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix}$  und  $\vec{b} = \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix}$  ist die Zahl

$$c = \vec{a} \cdot \vec{b} = |\vec{a}| \cdot |\vec{b}| \cdot \cos(\alpha)$$

Bei Gleichsetzen der beiden Berechnungsmethoden lässt sich leicht der Schnittwinkel berechnen. Dazu jedoch mehr in Kapitel??.

Eine Besonderheit des Skalarproduktes ist es, dass es 0 ist, wenn die Vektoren  $\vec{a}$  und  $\vec{b}$  senkrecht aufeinander stehen. Das Skalarprodukt kann also zur Überprüfung des berechneten Normalenvektors in Kapitel 17.1.10 dienen.

## 19. VEKTORPRODUKT

Das Vektorprodukt ist die zweite mögliche Art zwei Vektoren zu multiplizieren. Das Ergebnis ist, wie es sich aus dem Namen »Vektorprodukt« erkennen lässt, ein Vektor. Das Vektorprodukt der Vektoren  $\vec{a}$  und  $\vec{b}$  schreibt man  $\vec{v} = \vec{a} \times \vec{b}$ . Gelesen wird dies »a kreuz b«. Das Ergebnis eines Vektorproduktes ergibt sich mit

$$\vec{a} \times \vec{b} = \begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix} \times \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix} = \begin{pmatrix} a_2 \cdot b_3 - b_2 \cdot a_3 \\ a_3 \cdot b_1 - a_1 \cdot b_3 \\ a_1 \cdot b_2 - a_2 \cdot b_1 \end{pmatrix}$$

Der Ergebnisvektor steht senkrecht zu den beiden Vektoren, die miteinander multipliziert wurden. Daher eignet sich das Vektorprodukt, um einen Normalenvektor zu zwei gegebenen Vektoren zu bestimmen. Der Betrag eines Vektorproduktes ist die Fläche des von den beiden Vektoren  $\vec{a}$  und  $\vec{b}$  aufgespannten Parallelogramms. Also gilt  $A = |\vec{a} \times \vec{b}|$ . Die Fläche eines Dreiecks, das von den beiden Vektoren begrenzt ist, berechnet sich mit  $A = \frac{1}{2} |\vec{a} \times \vec{b}|$ , da dessen Fläche nur halb so groß wie die des Parallelogramms ist.

In  $\mathbb{R}^3$  gilt für den Betrag des Vektorproduktes  $|\vec{a} \times \vec{b}| = |\vec{a}| \cdot |\vec{b}| \cdot \sin(\alpha)$ . Für den Winkel zwischen den beiden Vektoren gilt folglich  $\sin(\alpha) = \frac{|\vec{a} \times \vec{b}|}{|\vec{a}| \cdot |\vec{b}|} = \frac{A_{\triangle}}{|\vec{a}| \cdot |\vec{b}|}$ . Bei gegebener Rechtecksfläche und bei gegebenen Vektoren  $\vec{a}$  und  $\vec{b}$  lässt sich mit obiger Formel der Winkel zwischen diesen beiden Vektoren berechnen.

Das Vektorprodukt kann auch im  $\mathbb{R}^2$  verwendet werden, um den Flächeninhalt eines Parallelogramms zu bestimmen. Man fügt einfach eine dritte Achse hinzu (d.h. eine dritte Komponente in allen Vektoren, die stets 0 ist). Aus  $\vec{a} = (a_1; a_2)^T$  und  $\vec{b} = (b_1; b_2)^T$  werden also  $\vec{a} = (a_1; a_2; 0)^T$  und  $\vec{b} = (b_1; b_2; 0)^T$ . Dann ist die Fläche des von  $\vec{a}$ ,  $\vec{b}$  im  $\mathbb{R}^2$  aufgespannten Parallelogramms:

$$\begin{aligned} A &= |\vec{a} \times \vec{b}| = \left| \begin{pmatrix} a_1 \\ a_2 \\ 0 \end{pmatrix} \times \begin{pmatrix} b_1 \\ b_2 \\ 0 \end{pmatrix} \right| \\ &= \left| \begin{pmatrix} 0 \\ 0 \\ a_1 b_2 - a_2 b_1 \end{pmatrix} \right| = \sqrt{(a_1 b_2 - a_2 b_1)^2} \\ &= |a_1 b_2 - a_2 b_1| = \left| \begin{vmatrix} a_1 & a_2 \\ b_1 & b_2 \end{vmatrix} \right| \end{aligned}$$

Das Vorzeichen der Determinante  $D = \begin{vmatrix} a_1 & a_2 \\ b_1 & b_2 \end{vmatrix}$  sagt dabei etwas über die Orientierung:

$D > 0$ :  $\vec{b}$  zeigt in die linke Halbebene von  $\vec{a}$

$D < 0$ :  $\vec{b}$  zeigt in die rechte Halbebene von  $\vec{a}$

**Example 62.** Vektorprodukt im  $\mathbb{R}^2$ , Orientierung von Vektoren, Parallelogrammfläche

**Example.** Gegeben sind

$$\vec{a} = \begin{pmatrix} 2 \\ 1 \end{pmatrix} \quad \vec{b} = \begin{pmatrix} -1 \\ -1 \end{pmatrix}$$

Dann ist

$$D = \begin{vmatrix} 2 & 1 \\ -1 & -1 \end{vmatrix} = -2 + 1 = -1 < 0$$

$$A = |D| = 1$$

$\vec{b}$  zeigt in die rechte Halbebene von  $\vec{a}$ , da  $D < 0$ .

### 19.1. Die Gesetze des Vektorproduktes.

- (1)  $\vec{a} \times \vec{a} = 0$
- (2)  $\vec{a} \times \vec{b} = 0 \Leftrightarrow \angle(\vec{a}, \vec{b}) = 0 \Leftrightarrow \vec{a} \parallel \vec{b}$
- (3)  $\vec{a} \times \vec{b} = -\vec{b} \times \vec{a}$
- (4)  $(r \cdot \vec{a}) \times (s \cdot \vec{b}) = r \cdot s \cdot (\vec{a} \times \vec{b})$
- (5)  $\vec{a} \times (\vec{b} + \vec{c}) = \vec{a} \times \vec{b} + \vec{a} \times \vec{c}$

## 20. SPATPRODUKT

**Definition 108.** Spatprodukt, Rechts- und Linkssystem erkennen

**Definition.** Das Spatprodukt  $[\vec{a}, \vec{b}, \vec{c}]$ <sup>12</sup> ist ...

- ... positiv (negativ), wenn  $\vec{a}, \vec{b}, \vec{c}$  ein Rechts-(Links-)System bilden.
- ... gleich 0, wenn  $\vec{a}, \vec{b}, \vec{c}$  komplanar sind, d.h. in einer Ebenen liegen.
- ... dem Betrag nach gleich dem Volumen des von  $\vec{a}, \vec{b}, \vec{c}$  aufgespannten Spats.

Möglichkeiten der Berechnung im  $\mathbb{R}^3$ :

$$\begin{aligned} [\vec{a}, \vec{b}, \vec{c}] &= (\vec{a} \times \vec{b}) \cdot \vec{c} = (\vec{b} \times \vec{c}) \cdot \vec{a} = (\vec{c} \times \vec{a}) \cdot \vec{b} \\ &= \begin{vmatrix} a_1 & a_2 & a_3 \\ b_1 & b_2 & b_3 \\ c_1 & c_2 & c_3 \end{vmatrix} \\ &= \begin{vmatrix} a_1 & b_1 & c_1 \\ a_2 & b_2 & c_2 \\ a_3 & b_3 & c_3 \end{vmatrix} \end{aligned}$$

Das Spatprodukt ist eine Mischung aus Skalarprodukt (siehe Kapitel 18) und aus Vektorprodukt (siehe Kapitel 19). Das Volumen des von  $\vec{a}, \vec{b}$  und  $\vec{c}$  aufgespannten Spats ist der Betrag des Spatproduktes:  $V_{\text{Spat}} = |(\vec{a} \times \vec{b}) \cdot \vec{c}|$ . Das Volumen einer dreiseitigen Pyramide lässt sich mit  $V_{\text{Pyramide}} = \frac{1}{6} |(\vec{a} \times \vec{b}) \cdot \vec{c}|$  berechnen.

Wenn die Vektoren  $\vec{a}, \vec{b}, \vec{c}$  linear abhängig sind, so gilt für das Spatprodukt  $(\vec{a} \times \vec{b}) \cdot \vec{c} = 0$ .

## 21. GERADEN

**21.1. Mittelpunkt einer Strecke.** Der Mittelpunkt einer Strecke berechnet sich folgendermaßen ( $A$  und  $B$  sind die Endpunkte der Strecke;  $\vec{a}$  und  $\vec{b}$  sind die Ortsvektoren dieser Punkte):  $\frac{1}{2} \cdot (\vec{a} + \vec{b})$ .

**21.2. Parameterdarstellung.** Es gibt zwei Möglichkeiten, die Gleichung einer Geraden zu schreiben. Eine davon ist die PARAMETERDARSTELLUNG.

Die Gerade in Parameterdarstellung besteht aus zwei Teilen:

- (1) *Der Stützvektor*  
Der *Stützvektor* gibt den Punkt an, von dem die Gerade beginnt, also der Punkt an den der Richtungsvektor »angelegt« wird.
- (2) *Der Richtungsvektor*  
Der *Richtungsvektor* gibt die Richtung in den Raum sowie die Ausrichtung (positives oder negatives Vorzeichen) an.

Eine Gerade schreibt man folgendermaßen:  $g: \vec{x} = \vec{x}_0 + r \cdot \vec{a}$ . Der Stützvektor ist hier  $\vec{x}_0$ , der Richtungsvektor ist  $\vec{a}$ . Der Faktor  $r$  heißt »Parameter«. Er gibt die »Verlängerung« der Richtungsvektor an.

**21.3. Normalenform.** Die Normalenform einer Geraden gibt eine Gerade an, die zur ursprünglichen Geraden senkrecht steht. Eine Geraden-Normalenform kann nur im  $\mathbb{R}^2$  existieren, da es in  $\mathbb{R}^3$  unendlich viele Geraden gibt, die auf einer anderen Geraden senkrecht stehen (also eine Ebene).

Als erstes berechnet man den Normalenvektor der Geraden. Wie dies funktioniert steht in 17.1.10. Dann multipliziert man beide Seiten der Geradengleichung mit diesem und erhält nun die Normalenform der Geraden.

Eine Normalenform der Geraden  $g: \vec{x} = \begin{pmatrix} 1 \\ 1 \end{pmatrix} + r \cdot \begin{pmatrix} 1 \\ 2 \end{pmatrix}$  ist zum Beispiel  $g: -2x + y = -1$ .

### 21.4. Geraden-Konstellationen.

<sup>12</sup>Alternative Schreibweise:  $(\vec{a}, \vec{b}, \vec{c})$ .

21.4.1. *Identische Geraden.* 2 Geraden sind identisch, wenn der Stützvektor der zweiten Geraden auf der ersten Geraden liegt und wenn die beiden Richtungsvektoren gleich sind. Dies lässt sich am besten an einem Beispiel veranschaulichen:

Es soll überprüft werden, ob die beiden Geraden  $g_1 : \vec{x} = \begin{pmatrix} 1 \\ -5 \end{pmatrix} + r \cdot \begin{pmatrix} 1 \\ 4 \end{pmatrix}$  und  $g_2 : \vec{x} = \begin{pmatrix} 2 \\ 6 \end{pmatrix} + r \cdot \begin{pmatrix} 3 \\ 2 \end{pmatrix}$  identisch sind oder nicht. Der Ortsvektor des Stützvektors der zweiten Geraden ist  $\begin{pmatrix} 3 \\ 2 \end{pmatrix}$ . Diesen setzt man mit der ersten Gerade gleich, um zu überprüfen, ob dieser auf dieser liegt:  $\begin{pmatrix} 3 \\ 2 \end{pmatrix} = \begin{pmatrix} 1 \\ -5 \end{pmatrix} + r \cdot \begin{pmatrix} 1 \\ 4 \end{pmatrix}$ . Dies ergibt  $r_1 = 1$  und  $r_2 = \frac{7}{4}$ . Da  $r_1$  und  $r_2$  nicht übereinstimmen, können beide Geraden nicht mehr identisch sein. Die zweite Überprüfung muss hier deshalb nicht mehr durchgeführt werden.

21.4.2. *Parallele Geraden.* Zwei Geraden sind parallel, wenn folgendes gilt:

- (1) Der Ortsvektor des Stützvektors der ersten Geraden liegt nicht auf der zweiten Geraden
- (2) Die Richtungsvektoren lassen sich so kürzen, dass sie gleich sind

Die Geraden  $g_1 : \vec{x} = \begin{pmatrix} 1 \\ -5 \end{pmatrix} + r \cdot \begin{pmatrix} 2 \\ 4 \end{pmatrix}$  und  $g_2 : \vec{x} = \begin{pmatrix} 2 \\ 6 \end{pmatrix} + r \cdot \begin{pmatrix} 4 \\ 8 \end{pmatrix}$  wären also parallel, da die beiden Bedingungen für sie erfüllt sind.

21.4.3. *Sich schneidende Geraden.* Geraden können sich in einem Punkt schneiden. Um zu sehen ob zwei Geraden dies tun und wenn ja, was der Schnittpunkt ist setzt man diese gleich. Für die Geraden  $g : \vec{x} = \vec{p} + r \cdot \vec{a}$  und  $h : \vec{x} = \vec{q} + s \cdot \vec{b}$  erhält man also die folgenden drei Gleichungen

$$\begin{aligned} 1.) \quad p_1 + r \cdot a_1 &= q_1 + s \cdot b_1 \\ 2.) \quad p_2 + r \cdot a_2 &= q_2 + s \cdot b_2 \\ 3.) \quad p_3 + r \cdot a_3 &= q_3 + s \cdot b_3 \end{aligned}$$

Mit Hilfe der Gleichungen 1.) und 2.) berechnet man Werte für  $r$  und  $s$ . Diese Werte setzt man zur Überprüfung in die Gleichung 3.) ein. Wenn diese erfüllt ist, existiert ein Schnittpunkt, ansonsten existiert kein Schnittpunkt zwischen den beiden Geraden.

Der Winkel zwischen den beiden Geraden berechnet sich als Winkel zwischen den beiden Richtungsvektoren:

$$\cos(\alpha) = \left| \frac{\vec{a} \cdot \vec{b}}{|\vec{a}| \cdot |\vec{b}|} \right|.$$

21.4.4. *Windschiefe Geraden.* Windschiefe Geraden können nur in einer Dimension existieren, die größer als 2 ist, also z.B. in  $\mathbb{R}^3$ . Windschiefe Geraden sind weder identisch oder parallel noch schneiden sie sich. Sie laufen also aneinander vorbei. Man kann auch sagen, dass beide Geraden in verschiedenen Ebenen liegen.

Um festzustellen ob zwei Geraden windschief sind prüft man zuerst, ob beide Geraden identisch oder parallel sind. Wenn sie dies nicht sind überprüft man beide auf einen gemeinsamen Schnittpunkt. Wenn dieser Punkt nicht existiert sind beide Geraden zueinander windschief.

21.5. **Die Hesse-Gleichung.** Die HESSE-GLEICHUNG ist eine besondere Art der Normalenform. An ihr kann man sofort den Abstand der Geraden zum Ursprung ablesen und mit ihr den Abstand zu anderen Punkten, Geraden oder Ebenen berechnen. Da die Normalenform einer Geraden nur in  $\mathbb{R}^2$  existieren kann, kann die Hesse-Gleichung auch nur in dieser Dimension existieren.

Die Grundform der HESSE-GLEICHUNG lautet:  $g : \vec{x} \cdot \vec{e}_n = \vec{x}_0 \cdot \vec{e}_n$ . Auf der rechten Seite kann man nun direkt den kleinsten Abstand der Geraden zum Ursprung ablesen. So ist zum Beispiel der Abstand der Geraden  $g : \vec{x} = \begin{pmatrix} 1 \\ 3 \end{pmatrix} + r \cdot \begin{pmatrix} 2 \\ 1 \end{pmatrix}$  zum Ursprung  $\sqrt{5}LE$ , da dies der rechte Teil der Hesse-Gleichung ist.

21.6. **Abstand zwischen Punkt und Gerade.** Die Berechnung des Abstandes zwischen einem Punkt und einer Geraden unterscheidet sich für  $\mathbb{R}^2$  und  $\mathbb{R}^3$ .

**In  $\mathbb{R}^2$ :** Der Term der Hesse-Gleichung (siehe Kapitel 21.5) lässt sich umformen nach  $g : (\vec{x} - \vec{x}_0) \cdot \vec{e}_n = 0$ .

Nun setzt man für  $\vec{x}$  den Ortsvektor des Punktes ein, zu dem man den Abstand bestimmen möchte.

Kurz gesagt ist die Gleichung des Abstandes  $d$  also folgende:  $d = |(\vec{x} - \vec{x}_0) \cdot \vec{e}_n|$ .

**In  $\mathbb{R}^3$ :** In  $\mathbb{R}^3$  ist die Berechnung des Abstandes schwieriger. Zur Veranschaulichung dient Abbildung 29.

Der Abstand beträgt  $d = |\overrightarrow{PF}|$ . Zu Berechnung ist es hilfreich zu wissen, dass  $\overrightarrow{PF} \perp \vec{a}$  gilt, also  $\overrightarrow{PF} \cdot \vec{a} = 0$  ergibt. Zuerst berechnet man also den Vektor  $\overrightarrow{PF}$ , multipliziert diesen dann mit  $\vec{a}$  um einen Wert für den Parameter der Geraden zu erhalten. Mit diesem lässt sich dann ein Zahlwert für  $\overrightarrow{PF}$  errechnen.

### 21.7. Abstand zwischen zwei Geraden.

**Definition 109.** Abstand zweier Geraden

**Definition.** Der Abstand  $d$  der beiden Geraden

$$\begin{aligned} g_1 : \vec{X} &= \vec{a} + t\vec{u} \\ g_2 : \vec{X} &= \vec{b} + t\vec{v} \end{aligned}$$

ist gegeben durch:

- falls  $\vec{u}, \vec{v}$  parallel:

$$d = \frac{|\left(\vec{b} - \vec{a}\right) \times \vec{u}|}{|\vec{u}|}$$

- falls  $\vec{u} \times \vec{v} \neq \vec{0}$ :

$$(21) \quad d = \frac{\left| \left[ \left(\vec{b} - \vec{a}\right), \vec{u}, \vec{v} \right] \right|}{|\vec{u} \times \vec{v}|}$$

21.7.1. *Identische Geraden.* Der Abstand zweier identischer Geraden beträgt immer 0. Siehe Abschnitt 21.4.1 zum Nachweis identischer Geraden.

21.7.2. *Parallele Geraden.* Die beiden Geradengleichungen lauten hier  $g_1 : \vec{x} = \vec{x}_0 + r \cdot \vec{a}$  und  $g_2 : \vec{x} = \vec{x}_1 + s \cdot \vec{b}$ . Zuerst muss nachgewiesen werden, dass besagte Geraden parallel sind (siehe 21.4.2). Dann kann der Abstand berechnet werden. Die Berechnung des Abstandes zweier paralleler Geraden in  $\mathbb{R}^2$  und  $\mathbb{R}^3$  ist unterschiedlich.

**In  $\mathbb{R}^2$ :** Man berechnet einfach den Abstand des Ortsvektors  $\vec{x}_1$  der zweiten Geraden zu der ersten Geraden. Wie dies geht steht in 21.6.

**In  $\mathbb{R}^3$ :** Hier verwendet man den Ortsvektor  $\vec{x}_1$  der zweiten Geraden für das Lotfußverfahren (siehe 21.6), um mit diesem den Abstand der ersten Geraden zur zweiten Geraden auszurechnen.

21.7.3. *Sich schneidende Geraden.* Zuerst muss nachgewiesen werden, dass die beiden Geraden sich schneiden (siehe Kapitel 21.4.3). Sich schneidende Geraden haben immer den Abstand 0, da sie einen Punkt gemeinsam haben.

21.7.4. *Windschiefe Geraden.* Gegeben sind die beiden zueinander windschiefen Geraden  $g : \vec{x} = \vec{x}_A + t \cdot \vec{a}$  und  $h : \vec{x} = \vec{x}_B + r \cdot \vec{b}$ . Um den Abstand dieser beiden Geraden zueinander auszurechnen stellt man die Ebene  $E : \vec{x} = \vec{x}_A + t \cdot \vec{a} + r \cdot \vec{b}$  auf. Diese Ebene enthält die erste Gerade und ist parallel zur zweiten Geraden, da sie beide Richtungsvektoren in sich vereint. Nun berechnet man einfach den Abstand  $d$  dieser Ebene zum Stützvektor  $\vec{b}$  (siehe Kapitel 22.3) der zweiten Geraden  $g$  und  $h$ . Dies ist gleichzeitig der Abstand der beiden Geraden.

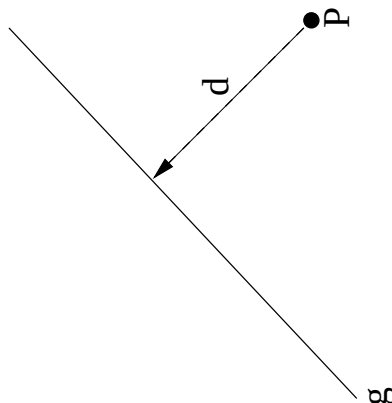


ABBILDUNG 29. Lotfußverfahren

**21.8. Durchstoßpunkte mit den Grundebenen.** Geraden schneiden die Grundebenen. Diese Schnittpunkte sollen berechnet werden. Es gibt im Normalfall drei Schnittpunkte die berechnet werden sollen:  $S_{12}$ ,  $S_{13}$  und  $S_{23}$ .

Wenn eine Geraden eine Ebene wie  $E: \vec{x} = r \cdot \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + s \cdot \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}$  schneidet, muss die  $x_3$  Koordinate 0 sein. Wir stellen also drei Gleichungen auf mit der Geraden  $g: \vec{x} = \vec{x}_P + t \cdot \vec{a}$  auf:

- 1.)  $p_1 + t \cdot a_1$
- 2.)  $p_2 + t \cdot a_2$
- 3.)  $p_3 + t \cdot a_3$

Da wir den Schnittpunkt  $S_{12}$  berechnen wollen, setzen wir die dritte Gleichung  $= 0$  um den Parameter  $t$  für diesen Schnittpunkt zu erhalten. Mit dem Parameter lassen sich die noch benötigten Koordinaten des Punktes berechnen.

## 22. EBENEN

**22.1. Schwerpunkt eines Dreiecks.** Der Schwerpunkt  $S$  eines Dreiecks berechnet sich ähnlich wie der Mittelpunkt einer Strecke (siehe 21.1). Für den Ortsvektor des Schwerpunktes gilt folgendes:

$$\vec{x}_S = \frac{1}{3}(\vec{x}_A + \vec{x}_B + \vec{x}_C)$$

Die Vektoren  $\vec{x}_A$ ,  $\vec{x}_B$  und  $\vec{x}_C$  sind die Ortsvektoren der Dreieckspunkte  $A$ ,  $B$  und  $C$ .

**22.2. Ebenengleichung.** Die Gleichung einer Ebene besteht aus drei Dingen:

- Dem Stützvektor
- Zwei Spannvektoren mit jeweils verschiedenen Parametern, die die Ebene vom Stützvektor aufbauen

Eine Ebene wird folgendermaßen geschrieben:  $E: \vec{x} = \vec{x}_A + r \cdot \vec{a} + s \cdot \vec{b}$ . Der Vektor  $\vec{x}_A$  ist der Stützvektor.  $\vec{a}$  und  $\vec{b}$  sind die Spannvektoren.

**22.2.1. Mit drei Punkten.** Die Gleichung einer Ebene kann mit Hilfe von 3 Punkten angegeben werden. Die aus den Punkten  $A$ ,  $B$  und  $C$  resultierende Ebene lautet  $E: \vec{x} = \vec{x}_A + r \cdot (\vec{x}_B - \vec{x}_A) + s \cdot (\vec{x}_C - \vec{x}_A)$ .

**22.2.2. Mit einem Punkt und zwei Spannvektoren.** Eine Ebene kann auch mit einem Punkt und zwei Spannvektoren angegeben werden. Die Ebene, die aus dem Punkt  $A$  und den Spannvektoren  $\vec{a}$  und  $\vec{b}$  besteht lautet  $E: \vec{x} = \vec{x}_A + r \cdot \vec{a} + s \cdot \vec{b}$ .

**22.2.3. Zwischen 2 Geraden.** Eine Ebene, die durch zwei Geraden angegeben wird ist leicht zu berechnen. Die resultierende Ebene besitzt als Punkt den Stützvektor der ersten Geraden und als zwei Spannvektoren die Richtungsvektoren der beiden Geraden. Die Ebene, die durch die beiden Geraden  $g: \vec{x} = \vec{x}_A + r \cdot \vec{a}$  und  $h: \vec{x} = \vec{x}_B + s \cdot \vec{b}$  lautet  $E: \vec{x} = \vec{x}_A + r \cdot \vec{a} + s \cdot \vec{b}$ .

**22.2.4. Normalenform der Ebene.** Der Normalenvektor  $\vec{n}$  einer Ebene  $E$  ist das Spatprodukt (siehe 20) des Einheitsvektors  $\vec{e}$  und der beiden Spannvektoren der Ebene. Der Normalenvektor der Ebene  $E: \vec{x} = \begin{pmatrix} 1 \\ 3 \\ 2 \end{pmatrix} +$

$$r \begin{pmatrix} -2 \\ 1 \\ 3 \end{pmatrix} + s \begin{pmatrix} 1 \\ 3 \\ 0 \end{pmatrix} \text{ lautet also } \vec{n} = \begin{pmatrix} -9 \\ 3 \\ -7 \end{pmatrix}.$$

**22.2.5. Hesseform der Ebene.** Die allgemeine Gleichung für die Hesseform einer Ebene lautet  $E: (\vec{x} - \vec{x}_0) \cdot \vec{e}_n = 0$  oder aber auch  $E: \vec{x} \cdot \vec{e}_n = \vec{x}_0 \cdot \vec{e}_n$ .

**22.3. Abstand zwischen Punkt und Ebene.** Der Abstand des Punktes wird mittels der Hesseform bestimmt. Man setzt den Punkt für den Vektor  $\vec{x}$  der ursprünglichen Hesseform ein und erhält somit — nachdem man den Betrag genommen hat — den Abstand. Der Abstand des Punktes  $P \begin{pmatrix} 1 \\ 3 \\ 2 \end{pmatrix}$  von der Ebene

$$E: \vec{x} = \begin{pmatrix} 4 \\ -2 \\ 2 \end{pmatrix} + r \begin{pmatrix} -2 \\ 1 \\ 3 \end{pmatrix} + s \begin{pmatrix} 1 \\ 3 \\ 0 \end{pmatrix} \text{ mit } \vec{n} = \begin{pmatrix} -9 \\ 3 \\ -7 \end{pmatrix} \text{ ist also } d = \left| \frac{-9+9-14+56}{\sqrt{139}} \right| = 3,56 \text{ LE.}$$



**22.4. Abstand zwischen Gerade und Ebene.** Der Abstand zwischen einer Geraden und einer Ebenen kann folgende Werte annehmen:

- $d = 0$ , wenn die Gerade mit der Ebene mindestens einen Punkt gemeinsam hat. Ein gemeinsamer Punkt existiert, wenn die Gerade in der Ebene enthalten ist, oder wenn die Gerade die Ebene schneidet.
- $d > 0$  wenn die Gerade mit der Ebene keinen Punkt gemeinsam hat. Dieser Fall kann nur eintreten, wenn die Gerade parallel zur Ebene steht.

Es gibt nun zwei Lösungsansätze:

- (1) Zuerst muss man also feststellen, ob ein gemeinsamer Punkt existiert. Falls dies nicht zutrifft muss man den Abstand der Geraden zur Ebene bestimmen.
- (2) Man berechnet  $\vec{n} \cdot \vec{a} = 0$ . Falls diese Gleichung wahr ist, ist die Gerade senkrecht zu E. Man muss also nun den Abstand berechnen, indem man den Stützvektor der Geraden in die Hesseform der Ebene einsetzt.

**22.5. Abstand zwischen zwei Ebenen.** Zwei Ebenen können entweder parallel sein oder sie scheiden sich. Im ersten Fall muss man den Abstand der einen Geraden zum Stützvektor der anderen berechnen. Im Falle des Schnittes beider Ebenen beträgt der Abstand automatisch 0.

**22.6. Schnitt zwischen Gerade und Ebene.** Eine Gerade kann natürlich eine Ebene schneiden. Um diesen zu berechnen, setzt man die  $x$ ,  $y$  und  $z$  Werte der gesamten Geradengleichung für die  $x$ ,  $y$  und  $z$  Werte der Ebenen-Normalenform ein. Diese Berechnung wird erfolgreich sein, wenn ein Schnittpunkt existiert. Ansonsten geht die Gleichung nicht auf.

Im Falle eines existenten Schnittpunkts ist das Resultat der obigen Rechnung ein Wert für den Parameter der Geradengleichung. Mit diesem Wert kann man nun die Koordinaten des Schnittpunkts berechnen.

**22.6.1. Schnittwinkel.** Den Schnittwinkel zwischen einer Geraden und einer Ebene berechnet mit Hilfe des Richtungsvektors der Geraden und mit dem Normalenvektor der Ebene. Die Formel lautet:  $\sin(\alpha) = \left| \frac{\vec{n} \cdot \vec{a}}{|\vec{n}| \cdot |\vec{a}|} \right|$ .

**22.7. Schnitt zweier Ebenen.** Zwei Ebenen können sich schneiden. Das entstehende Schnittgebilde ist eine Gerade. Eine Eigenschaft dieser Schnittgeraden ist es, dass ihr Richtungsvektor senkrecht zu den Normalenvektoren beider Ebenen steht. Man berechnet den Richtungsvektor der Schnittgeraden also, indem man mit Hilfe des Spatproduktes einen Vektor  $\vec{a}$  berechnet, der zu den Normalenvektoren  $\vec{n}_1$  und  $\vec{n}_2$  der Ebenen senkrecht steht.

Nun fehlt nur noch der Stützvektor der Geraden. Man setzt hierzu bei beiden Ebenen die  $z$  Koordinate 0 um damit ein Gleichungssystem von 2 Gleichungen mit 2 Unbekannten zu bekommen. Die erhaltenen  $x$  und  $y$  Werte sind die Koordinaten eines Punktes, der in beiden Ebenen liegt. Dies ist der Stützvektor.

**22.7.1. Schnittwinkel.** Der Schnittwinkel zwischen zwei sich schneidenden Ebenen rechnet man mit Hilfe ihrer beiden Normalenvektoren aus. Die Formel lautet:  $\cos(\alpha) = \left| \frac{\vec{n}_1 \cdot \vec{n}_2}{|\vec{n}_1| \cdot |\vec{n}_2|} \right|$ . Wenn  $\alpha = 0^\circ$  gilt, dann sind die beiden Ebenen parallel.

**22.8. Spurpunkte.** Spurpunkte einer Ebene ähneln den Spurpunkten einer Geraden (siehe 21.8).

Man rechnet am besten zu Beginn die Normalenform der Ebene aus. Im Gegensatz zur Berechnung der Spurpunkte bei Geraden setzt man hier zwei Koordinaten = 0, da ja zwei Parameter berechnet werden müssen.

- Um den Spurpunkt  $A_1$  zu erhalten setzt man  $x_2 = x_3 = 0$
- Um den Spurpunkt  $A_2$  zu erhalten setzt man  $x_1 = x_3 = 0$
- Um den Spurpunkt  $A_3$  zu erhalten setzt man  $x_1 = x_2 = 0$

**22.9. Spurgeraden.** Es gibt drei Spurgeraden, die man ermitteln muss:  $S_{12}$ ,  $S_{13}$  und  $S_{23}$ . Die Spurgeraden sind die Verbindungsgeraden zwischen den Spurpunkten  $A_1$ ,  $A_2$  und  $A_3$  der Ebene.

## 23. SPIEGELUNGEN

### 23.1. Punkte.

**23.1.1. An einem Punkt.** Der Punkt  $P$  soll an dem Punkt  $Q$  gespiegelt werden. Es gilt für den entstehenden Spiegelpunkt  $\vec{x}_{P'} = \vec{x}_P + 2 \cdot \vec{PQ}$ . Um das Ergebnis zu überprüfen, muss die Gleichung  $\vec{x}_Q = \vec{x}_P + \frac{1}{2} \cdot \vec{PP'}$  erfüllt sein. Der Punkt  $Q$  muss also die Mitte der Strecke  $PP'$  sein.

**23.1.2. An einer Geraden in  $\mathbb{R}^2$ .** Der Punkt  $P$  soll an der Geraden  $g : \vec{x} \cdot \vec{e}_n = \vec{x}_0 \cdot \vec{e}_n$  gespiegelt werden. Die Formel dafür lautet  $\vec{x}_{P'} = \vec{x}_P - 2 \cdot d \cdot \vec{e}_n$ . Der Abstand  $d$  wird mit dem Vorzeichen eingesetzt, um den richtigen Spiegelpunkt zu bekommen. Der Vektor  $\vec{e}_n$  ist der Normalenvektor der Geraden  $g$ .

Um das Ergebnis auf Richtigkeit zu überprüfen setzt man den Mittelpunkt der Strecke  $\overline{PP'}$  in die Gerade  $g$  ein. Wenn diese Gleichung erfüllt ist, ist der Punkt  $P'$  der Spiegelpunkt von  $P$ .

23.1.3. *An einer Geraden in  $\mathbb{R}^3$ .* Der Punkt  $P$  soll in  $\mathbb{R}^3$  an der Geraden  $g$  auf seinen Spiegelpunkt  $P'$  gespiegelt werden.

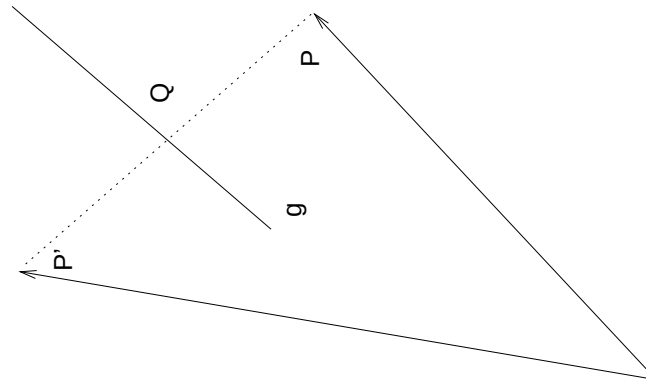


ABBILDUNG 30. Spiegelung Punkt an Gerade in  $\mathbb{R}^3$

Zur Veranschaulichung dient Abb. 30. Für diese Spiegelung gilt  $\vec{x}_{P'} = \vec{x}_P + \overrightarrow{PP'}$ . Es gilt weiterhin  $\overrightarrow{PP'} = 2 \cdot \overrightarrow{PQ}$ . Den Punkt  $Q$  kann man mit Hilfe des Lotfußverfahrens (siehe 21.6) berechnen. Wenn  $Q$  bekannt ist, lässt sich die obige Gleichung zu  $\vec{x}_{P'} = \vec{x}_Q + \overrightarrow{PQ}$  vereinfachen. Wenn man sein Ergebnis überprüfen will, berechnet man den Mittelpunkt der Strecke  $\overline{PP'}$  und setzt diesen in die Gerade  $g$  ein. Wenn diese Gleichung mit dem errechneten Punkt erfüllt ist, ist der Spiegelpunkt  $P'$  korrekt.

23.1.4. *An einer Ebene.* Nun soll der Punkt  $P$  an der Ebene  $E$  gespiegelt werden. Zuerst überprüft man, ob der Punkt  $P$  in der Ebene liegt. Wenn dies zutrifft, ist er mit seinem Spiegelpunkt identisch.

Wenn er nicht in der Ebene liegt gilt  $\vec{x}_{P'} = \vec{x}_P - 2 \cdot d \cdot \vec{e}_n$ , wobei  $\vec{e}_n$  der Einheitsnormalenvektor der Ebene ist.

## 23.2. Geraden.

23.2.1. *An einer Geraden.* Die Berechnung ist unterschiedlich, wenn zwei parallele, zwei windschiefe oder wenn zwei sich schneidende Geraden aneinander gespiegelt werden sollen.

Wir gehen davon aus, dass die Gerade  $h : \vec{x} = \vec{x}_A + t \cdot \vec{a}$  an der Geraden  $g : \vec{x} = \vec{x}_B + s \cdot \vec{b}$  gespiegelt werden soll.

**Zwei parallele Geraden:** Der Punkt  $A$  der Geraden  $h$  muss an der Geraden  $g$  gespiegelt werden (Punkt an Gerade spiegeln siehe 23.1.3). Der Richtungsvektor der gespiegelten Gerade ist der selbe wie in der ursprünglichen Geraden  $h$ . Die gespiegelte Gerade lautet also  $h' : \vec{x} = \vec{x}_{A'} + t \cdot \vec{a}$ .

**Zwei windschiefe Geraden:** Man stelle eine Ebene  $E$  mit Hilfe der Geraden  $g$  und dem Richtungsvektor der Geraden  $h$  auf. Jetzt kann man den Stützvektor von  $h$  an der Ebene  $E$  spiegeln (siehe 23.1.4). Die Spiegelgerade lautet  $h' : \vec{x} = \vec{x}_{A'} + t \cdot \vec{a}$ , der Richtungsvektor bleibt also gleich.

**Zwei schneidende Geraden:** Man berechnet zuerst den Schnittpunkt  $S$  der beiden Geraden. Dieser Punkt ist auch in der Spiegelgeraden enthalten. Jetzt benötigt man noch einen weiteren Punkt, der gespiegelt werden soll. Dies ist der Punkt  $A$  der Geraden  $h$ . Die Spiegelgerade lautet nun also  $h' : \vec{x} = \vec{x}_S + t \cdot (\vec{SA}')$ .

23.2.2. *An einer Ebene.* Gegeben sind die Ebene  $E : \vec{x} \cdot \vec{n} = \vec{x}_0 \cdot \vec{n}$  und die Gerade  $g : \vec{x} = \vec{x}_1 + t \cdot \vec{a}$ . Zuerst muss überprüft werden, ob die Gerade zur Ebene parallel ist (mit Hilfe von  $\vec{a} \cdot \vec{n} = 0$ ). Wenn dies zutrifft, muss überprüft werden, ob die Gerade in der Ebene oder außerhalb von dieser liegt. Wenn sie in dieser enthalten ist, ist sie identisch mit ihrer Spiegelgeraden. Wenn sie zur Ebene parallel liegt, muss nur der Stützvektor gespiegelt werden, um die Spiegelgerade zu erhalten.

Wenn die Gerade nicht zur Ebene parallel liegt, muss sie diese logischerweise schneiden. In diesem Fall ist die einfachste Möglichkeit die Gerade zu spiegeln, den Schnittpunkt mit der Ebene zu berechnen und anschließend einen weiteren Punkt von  $g$  an  $E$  zu spiegeln. Mit zwei Punkten der Spiegelgeraden ist ohne weiteres möglich, diese mit der zwei Punkte Form aufzustellen.

## 23.3. Ebenen.

23.3.1. *An einer Ebenen.* Die Ebene  $F : \vec{x} \cdot \vec{n}_F = \vec{x}_0 \cdot \vec{n}_F$  soll an der Ebene  $E : \vec{x} \cdot \vec{n}_E = \vec{x}_0 \cdot \vec{n}_E$  gespiegelt werden. Zuerst muss überprüft werden, ob die beiden Ebenen zueinander parallel liegen. Wenn dies zutrifft, muss überprüft werden, ob beide identisch oder parallel liegen. Wenn beide identisch sind, muss nichts getan werden. Wenn beide parallel liegen, spiegelt man einen Punkt  $A$  (am besten den Durchstoßpunkt) von  $F$  an der Ebene  $E$  auf den Punkt  $A'$ . Die Gleichung der Spiegelebene lautet nun  $F' : \vec{x} \cdot \vec{n}_F = \vec{x}_{A'} \cdot \vec{n}_F$ .

Wenn die beiden Ebenen weder identisch noch parallel sind, müssen sie sich logischerweise schneiden. Man berechnet nun die Schnittgerade, da sie sowohl in der Ebene  $F$  als auch in  $F'$  enthalten ist. Nun spiegelt man einen Punkt  $B$  der Ebene  $E$  auf den Punkt  $B'$ . Die Gleichung der Spiegelebene lautet nun  $F : \vec{x} = \vec{x}_s + t \cdot \vec{a} + r \cdot (\vec{x}_{B'} - \vec{x}_s)$ . Man ergänzt zur Schnittgerade also noch einen weiteren Spannvektoren, um die komplette Ebene zu erhalten.

## 24. KREISE UND KUGELN

24.1. **Kreis- und Kugelgleichung.** Kreis- und Kugelgleichungen sind ähnlich. Beide haben einen Mittelpunkt und einen Radius. Die allgemeine Gleichung für Kugel und Kreis lautet  $K : (\vec{x} - \vec{x}_M)^2 = r^2$ . Für einen Kreis hat der Vektor  $\vec{x}$  die Dimension 2, für eine Kugel die Dimension 3. Der Vektor  $\vec{x}_M$  gibt den Mittelpunkt des Kreises bzw. der Kugel an.

### 24.2. Kreis.

24.2.1. *Schnitt mit Geraden.* Ein Kreis kann von einer Geraden geschnitten werden. Die Gerade kann Passante (kein Schnittpunkt), Tangente (1 Schnittpunkt) oder Sekante (2 Schnittpunkte) sein. Um das Schnittgebilde zu errechnen, setzt man einfach die Gerade in die Kugelgleichung ein. Wenn man als Lösung keinen Wert erhält, existiert kein Schnittgebilde. Bei einer Lösung existiert ein einzelner Schnittpunkt, bei zwei Lösungen existieren zwei Schnittpunkte.

### 24.3. Kugel.

24.3.1. *Schnitt mit Geraden.* Um das Schnittgebilde zwischen einer Kugel und einer Geraden zu berechnen, verfährt man wie bei dem Schnitt zwischen Gerade und Kreis (siehe 24.2.1).

24.3.2. *Schnitt mit Ebene.* Eine Kugel kann von einer Ebene geschnitten werden. Bei vorhandenem Schnittgebilde können ein Schnittpunkt oder ein Schnittkreis entstehen. Um diese herauszubekommen, ermittelt man den Abstand  $d$  des Kugelmittelpunktes zur Ebene  $E$ . Gilt  $r = d$ , dann existiert nur ein Schnittpunkt (d.h. ein Schnittkreis mit dem Radius 0). In diesem Fall muss der Berührungspunkt  $\vec{x}_B = \vec{x}_M + d \cdot \vec{e}_n$  berechnet werden.

Gilt jedoch  $d < r$ , dann existiert ein Schnittkreis. Dieser hat den Mittelpunkt  $\vec{x}_{M_1} = \vec{x}_M + d \cdot \vec{e}_n$  und den Radius  $r' = \sqrt{r^2 - d^2}$ .

24.3.3. *Schnitt mit Kugel.* Eine Kugel kann wiederum eine Kugel schneiden. Um zu prüfen, ob die Kugeln sich schneiden, rechnet man den Abstand der beiden Kugelmittelpunkte  $P_1$  und  $P_2$  aus. Gilt  $|\vec{P_1 P_2}| = r_1 + r_2$ , dann berühren sich die beiden Kugeln in dem Punkt  $\vec{x}_B = \vec{x}_{M_1} + r_1 \cdot \vec{e}_{PP'}$ .

Gilt  $|\vec{P_1 P_2}| < r_1 + r_2$ , dann entsteht ein Schnittkreis mit dem Mittelpunkt  $\vec{x}_{M_1} = \vec{x}_M + d \cdot \vec{e}_n$  und dem Radius  $r' = \sqrt{r^2 - d^2}$ .

24.4. **Tangenten und Tangentialebenen.** Die Gleichung, um eine Tangente an einen Kreis bzw. eine Tangentialebene an eine Kugel aufzustellen lautet  $(\vec{x} - \vec{x}_M) \cdot (\vec{x}_B - \vec{x}_M) = r^2$ . Die Vektoren  $\vec{x}_M$  und  $\vec{x}_B$  sind die Ortsvektoren des Mittelpunktes  $M$  und des Berührungspunktes  $B$ .

Um in der Probe zu überprüfen, ob die Koordinaten des Berührungspunktes richtig sind, setzt man diesen in die Ebene und in die Kugel ein. Wenn die beiden Gleichungen erfüllt sind, sind die Koordinaten des Berührungspunktes korrekt.

## 25. FOLGEN UND REIHEN

In vielen Fragestellungen interessiert das Verhalten der  $a_i$  für  $i \rightarrow \infty$ .

**Definition 110.** Nullfolge

**Definition.**  $(a_n)$  ist eine Nullfolge, wenn  $|a_n|$  beliebig klein wird für hinreichend großes  $n$ .

**Definition 111.** Konvergente Folge (Folge mit Grenzwert  $A$ )

**Definition.**  $(a_n)$  konvergiert gegen  $A \in \mathbb{R}$ , wenn die Folge  $(a_n - A)$  eine Nullfolge ist. Sprechweise: Der Grenzwert der Folge  $(a_n)$  ist  $A$ . Schreibweise:  $\lim_{n \rightarrow \infty} a_n = A$ .

**Example 63.** Konvergente Folge

**Example.** Sei eine Folge  $a_i = \frac{3i^2-1}{i^2+1}$ . Ihr Grenzwert ist:

$$\begin{aligned}\lim_{i \rightarrow \infty} a_i &= \lim_{i \rightarrow \infty} \frac{3i^2 - 1}{i^2 + 1} \\ &= \lim_{i \rightarrow \infty} \frac{i^2 \left(3 - \frac{1}{i^2}\right)}{i^2 \left(1 + \frac{1}{i^2}\right)} \\ &= 3\end{aligned}$$

Das heißt: Der Abstand von  $a_i$  von 3 wird beliebig klein für  $i \rightarrow \infty$ . Ab welchem  $i$  gilt zum Beispiel  $|a_i - 3| < \frac{1}{10000}$ ?

$$\begin{aligned}|a_i - 3| &\stackrel{!}{<} \frac{1}{10000} \\ \Rightarrow \left| \frac{3i^2 - 1}{i^2 + 1} - 3 \right| &\stackrel{!}{<} \frac{1}{10000} \\ \Leftrightarrow \left| \frac{-4}{i^2 + 1} \right| &\stackrel{!}{<} \frac{1}{10000} \\ \Leftrightarrow |-4| \cdot 10000 &\stackrel{!}{<} i^2 + 1 \\ \Rightarrow 39999 &\stackrel{!}{<} i^2 \\ \Rightarrow i &> \sqrt{39999}\end{aligned}$$

Eine solche Folge  $(a_i)$ , für die die Bedingung  $|a_i - 3| < \frac{1}{10000}$  gelten soll, kann man also etwa bei  $i_0 = 200$  beginnen lassen.

**Definition 112.** Bestimmt divergente Folge

**Definition.**  $(a_i)$  heißt bestimmt divergent gegen  $+\infty$  [ $-\infty$ ], in Zeichen  $\lim_{i \rightarrow \infty} a_i = \infty$  [ $\lim_{i \rightarrow \infty} a_i = -\infty$ ], wenn zu jedem  $K > 0$  [ $K < 0$ ] ein  $i_0 = i_0(K)$  existiert mit  $a_i > K$  [ $a_i < K$ ] für alle  $i \geq i_0$ . Achtung:  $+\infty$  und  $-\infty$  sind keine Grenzwerte; sie werden gelegentlich als uneigentliche Grenzwerte bezeichnet.

**Example 64.** Divergente und bestimmt divergente Folge

**Example.**

- Die Folge  $a_i = 5^i$  ist bestimmt divergent gegen  $+\infty$ , d.h.  $\lim_{i \rightarrow \infty} a_i = \infty$ .
- Die Folge  $a_i = (-5)^i$  ist divergent<sup>13</sup>.

Die Monotonität einer Folge, zur Unterscheidung divergenter und bestimmt divergenter Folgen, prüft man durch Untersuchung von  $a_i - a_{i-1}$  oder  $\frac{a_i}{a_{i-1}}$ .

## 26. DIFFERENTIALRECHNUNG

**26.1. Symmetrie.** Eine Funktion kann entweder achsensymmetrisch oder punktsymmetrisch sein. Hier wird im Normalfall nur die Achsensymmetrie zur y-Achse und die Punktsymmetrie zum Ursprung überprüft.

**26.1.1. Punktsymmetrie zum Ursprung.** Hier muss die Beziehung  $f(x) = -f(-x)$  gelten. Die Werte auf positiver und negativer Seite unterscheiden sich also nur durch ihr Vorzeichen.

**26.1.2. Achsensymmetrie zur y-Achse.** Hier muss die Beziehung  $f(x) = f(-x)$  gelten. Die Funktionswerte sind also sowohl auf positiver als auch auf der negativen Seite gleich.

**26.1.3. Zu einem beliebigen Punkt.** Es muss die Beziehung  $f(x_0 - h) + f(x_0 + h) = 2 \cdot f(x_0)$  erfüllt sein, wenn die Funktion  $f(x)$  zum Punkt  $P(x_0 | f(x_0))$  symmetrisch sein soll.

**26.1.4. Zu einer beliebigen Achse.** Hier muss die Beziehung  $f(x_0 + h) = f(x_0 - h)$  erfüllt sein, wenn die Funktion  $f(x)$  zur Achse durch  $x_0$  symmetrisch sein soll.

**Example 65.** Ist die Funktion  $f(x) = \frac{1}{x^2 - 2x + 2}$  achsensymmetrisch durch  $x_0 = 1$ ?

**Example.** Es muss  $f(x_0 + h) = f(x_0 - h)$  gelten. Wir erhalten

$$f(1 + h) = \frac{1}{(1 + h)^2 - 2(1 + h) + 2} = \frac{1}{1 + 2h + h^2 - 2 - 2h + 2} = \frac{1}{h^2 + 1}$$

und

$$f(1 - h) = \frac{1}{(1 - h)^2 - 2(1 - h) + 2} = \frac{1}{1 - 2h + h^2 - 2 + 2h + 2} = \frac{1}{h^2 + 1}$$

Da die beiden Funktionswerte übereinstimmen, ist die Funktion achsensymmetrisch zu  $x_0 = 1$ .

<sup>13</sup>sie »geht gegen  $+\infty$  und  $-\infty$ «

**26.2. Verhalten.** Es gibt Funktionen, die gegen einen bestimmten Funktionswert  $y_0$  laufen, wenn die  $x$  Werte auch gegen einen bestimmten Wert  $x_0$  gehen.  $x_0$  kann Werte wie  $x_0 = +\infty$ ,  $x_0 = -\infty$  oder  $x_0 = 5$  annehmen. Es soll nun untersucht werden, ob die Funktion an der Stelle  $x_0$  gegen einen bestimmten Grenzwert läuft oder nicht. Wenn ja, soll angegeben werden, was dieser Wert ist.

Um das Verhalten für die Stelle  $x_0$  zu erforschen, berechnet man  $\lim_{x \rightarrow x_0} f(x)$ . Falls ein Wert existiert, ist dies der gesuchte Grenzwert  $y_0$ .

**26.2.1. Satz von l'Hopital.** Der Satz von l'Hopital dient zur Berechnung der Grenzwerte für  $x \rightarrow \infty$ . Nach dem Satz von l'Hopital bleibt der Grenzwert gleich, wenn man den Zähler und den Nenner einzeln ableitet. Dies vereinfacht das ganze wesentlich.

**Example 66.** Verhalten der Funktion  $f(x) = \frac{2+3e^x}{e^x}$  für  $x_0 = +\infty$

**Example.** Wir müssen  $\lim_{x \rightarrow x_0} f(x)$  berechnen. Es gilt nun also

$$\begin{aligned} & \lim_{x \rightarrow x_0} \left( \frac{2+3e^x}{e^x} \right) \quad | \text{l'Hopital} \\ &= \lim_{x \rightarrow x_0} \left( \frac{3e^x}{e^x} \right) \quad | \text{gekürzt} \\ &= \lim_{x \rightarrow x_0} \left( \frac{3}{1} \right) \\ &= 3 \end{aligned}$$

Unser gesuchter Grenzwert ist also  $y_0 = 3$ . Wir können nun festhalten, dass die Funktion für  $x \rightarrow +\infty$  gegen  $y = 3$  läuft.

**26.3. Polstellen.** Polstellen<sup>14</sup> treten nur bei gebrochenrationalen Funktionen auf. Sie liegen an den Stellen, wo der Nenner 0 wird, d.h. ein Ergebnis unmöglich ist, weil durch 0 geteilt wird.

**26.3.1. Polstelle mit Vorzeichenwechsel.** Eine Nullstelle des Nennerpolynoms, die nicht gleichzeitig Nullstelle des Zählerpolynoms ist und *ungerade Vielfachheit* besitzt (z.B. dreifache Nullstelle), ist eine Polstelle *mit VZW* (Vorzeichenwechsel). Sie wird durch eine senkrechte Asymptote gekennzeichnet, an die sich der Graph nähert.

**26.3.2. Polstelle ohne Vorzeichenwechsel.** Eine Nullstelle des Nennerpolynoms, die nicht gleichzeitig Nullstelle des Zählerpolynoms ist und *gerade Vielfachheit* besitzt (z.B. zweifache Nullstelle), ist eine Polstelle *ohne VZW* (Vorzeichenwechsel). Sie wird durch eine senkrechte Asymptote gekennzeichnet, an die sich der Graph nähert.

**26.3.3. Definitionslücken.**

**Definition 113.** Definitionslücke

**Definition.**  $x_0$  heißt (Definitions-)Lücke der Art  $\frac{0}{0}$  oder Unstetigkeitsstelle der Art  $\frac{0}{0}$ , wenn es eine Nullstelle von Zähler- und Nennerpolynom ist.

An der Stelle einer Definitionslücke besteht keine senkrechte Asymptote, sondern nur eine Lücke, die im Graphen durch ein unausgefülltes kleines Quadrat gekennzeichnet wird. Mehrfache Definitionslücken werden nur einfach angegeben.

**Definition 114.** Behebbarer Definitionslücke

**Definition.** Eine Definitionslücke  $x_0$  der Art  $\frac{0}{0}$  heißt »behebbarer Lücke«, wenn  $f(x)$  so zu  $f^*(x)$  gekürzt werden kann, dass  $x_0$  keine Definitionslücke von  $f^*(x)$  ist<sup>15</sup>. Die Lücke heißt dann »mit  $f^*(x_0)$  behebbar«.

**Example 67.** Behebbarer Definitionslücke

**Example.** Die Funktion  $y = \frac{x^2-1}{x+1}$  ist für  $x_0 = -1$  nicht definiert. Es gilt für alle  $x \neq x_0$ :

$$f(x) = \frac{x^2-1}{x+1} = \frac{(x+1)(x-1)}{x+1} = x-1 = f^*(x)$$

Man sagt: Die Lücke ist mit  $f^*(-1) = -2$  behebbar.

**Example 68.** Nicht behebbarer Definitionslücke von  $f(x) = \frac{(x-1)^2}{\ln(x)}$ ;  $\mathbb{D} = \mathbb{R}_0^+$

**Example.** Die Nullstellen des Zählers sind  $x_0 = 1$ . Die Nullstellen des Nenners sind  $x_1 = 1$ . An der Stelle  $x_0 = 1$  ist also eine Definitionslücke. Sie ist nicht behebbar.

**26.4. Asymptoten.**

**26.4.1. Senkrechte Asymptoten.** Senkrechte Asymptoten liegen an Polstellen (siehe Kapitel 26.3) und treten daher nur bei gebrochenrationalen Funktionen auf.

<sup>14</sup> auch: »Unendlichkeitsstellen«

<sup>15</sup>  $f^*(x)$  hat hier also eine Nullstelle, Polstelle oder einen Funktionswert

26.4.2. *Nicht-Senkrechte Asymptoten.* Anders als bei ganzrationalen Funktionen können sich gebrochenrationale Funktionen für  $x \rightarrow \pm\infty$  einer bestimmten Geraden oder einer anderen Kurve annähern. Eine solche Gerade heißt Asymptote. Eine Asymptote gibt nur die Richtung für sehr große  $|x|$  an, daher darf sie im Bereich um 0 von dem Graphen geschnitten werden.

Für die gebrochenrationale Funktion  $f(x) = \frac{a_0 + a_1x + \dots + a_nx^n}{b_0 + b_1x + \dots + b_mx^m}$  ( $a_n \neq 0$ ;  $b_m \neq 0$ ) können verschiedene Asymptoten existieren.

(1)  $n < m$

Die x-Achse ist waagerechte Asymptote, da der Nenner schneller wächst als der Zähler.

(2)  $n = m$

Die Gerade mit der Gleichung  $f(x) = \frac{a_n}{b_m}$  ist hier waagerechte Asymptote.

(3)  $n = m + 1$

$n$  ist hier also um 1 größer als  $m$ . Es entsteht eine schiefe Asymptote, die der ganzzahlige Teil der Polynomdivision  $\frac{\text{Zähler}}{\text{Nenner}}$  ist.

(4)  $n > m + 1$

$n$  ist hier also um mehr als 1 größer als  $m$ . Nun ist eine Kurve Asymptote; ihre Gleichung ist der ganzzahlige Teil der Polynomdivision  $\frac{\text{Zähler}}{\text{Nenner}}$  ist.

26.5. **Nullstellen.** Nullstellen einer Funktion liegen an genau den Stellen, an denen die Funktionswerte der betrachteten Funktion 0 sind. Es muss also die Bedingung  $f(x) = 0$  gelten. Die Lösungen dieser Gleichung sind die x Werte der Nullstellen. Der y-Wert an diesen Stellen ist natürlich immer 0. Wenn man vorher die Symmetrie berechnet, kann man sich bei vorhandener Symmetrie unter Umständen einige Arbeit sparen.

**Example 69.** Nullstellen der Funktion  $f(x) = x^2 - x - 3$

**Example.** Es muss gelten  $x^2 - x - 3 = 0$ . Es gilt jetzt:

$$\begin{aligned} x^2 - x &= 3 \\ \Leftrightarrow (x - \tfrac{1}{2})^2 &= 3\frac{1}{4} \\ \Leftrightarrow x_1 &= -\sqrt{3\frac{1}{4}} + \frac{1}{2} \\ \vee x_2 &= \sqrt{3\frac{1}{4}} + \frac{1}{2} \end{aligned}$$

Die Nullstellen der Funktion  $f(x) = x^2 - x - 3$  sind also  $N_1(-\sqrt{3\frac{1}{4}} + \frac{1}{2} \mid 0)$  und  $N_2(\sqrt{3\frac{1}{4}} + \frac{1}{2} \mid 0)$ .

26.6. **Ableitungen.** Ableiten bedeutet, dass man eine Funktion  $f(x)$  so umrechnet, dass man eine Funktion erhält, die die Steigung der Funktion  $f(x)$  an jeder beliebigen Stelle  $x$  angibt. Die berechnete Funktion wird »Ableitungsfunktion  $f'(x)$ « genannt.

Es steht fest, dass konstante Faktoren (z.B. Zahlen oder Parameter) durch die Ableitung nicht verändert werden. Sie bleiben einfach stehen und tauchen in der Ableitung unverändert auf.

26.6.1. *Ableiten von Potenzen.* Potenzen wie  $f(x) = x^3$  werden abgeleitet, indem man die Potenz so wie sie ist vor den Term schreibt und anschließend die Hochzahl um 1 verringert. Die Ableitung würde hier also  $f'(x) = 3x^2$  lauten. Die allgemeine Schreibweise der Ableitung der Funktion  $f(x) = x^a$  ist  $f'(x) = a \cdot x^{a-1}$ .

Die Ableitung der Funktion  $f(x) = x$  ist folglich  $f'(x) = 1$ . Die nächste Ableitung dieser Funktion ist  $f''(x) = 0$ . Die Ableitung einer Konstanten ist immer 0. Bildlich gesehen ist die Steigung einer zur x-Achse parallelen Geraden immer 0.

**Example 70.** Ableitung von  $f(x) = \sqrt{x}$

**Example.** Es gilt  $f(x) = \sqrt{x} = x^{\frac{1}{2}}$ . Also ist die Ableitung  $f'(x) = \frac{1}{2} \cdot x^{-\frac{1}{2}} = -\frac{1}{2\sqrt{x}}$ .

26.6.2. *Ableiten von Produkten — die Produktregel.* Die Produktregel der Differentialrechnung besagt: »Die Ableitung der Funktion  $f(x) = u(x) \cdot v(x)$  ist  $f'(x) = u'(x) \cdot v(x) + u(x) \cdot v'(x)$ «. Man sucht sich also die beiden Produkte, die jeweils die Variable  $x$  enthalten (also keine konstanten Faktoren) und leitet diese nach der obigen Regel ab.

**Example 71.** Ableitung der Funktion  $f(x) = x^3 \cdot e^x$  ab.

**Example.** Wie suchen uns als erstes die beiden Teilfunktionen  $u$  und  $v$ . Diese sind hier  $u(x) = x^3$  und  $v(x) = e^x$ . Damit können wir nun die Ableitung errechnen.

$$\begin{aligned} f'(x) &= u'(x) \cdot v(x) + u(x) \cdot v'(x) \\ &= 3x^2 \cdot e^x + x^3 \cdot e^x \\ &= e^x \cdot (3x^2 + x^3) \\ &= e^x \cdot x^2(x + 3) \end{aligned}$$

26.6.3. *Ableiten von Brüchen — die Quotientenregel.* Das Ableiten von Quotienten ist ähnlich wie das Ableiten von Produkten. Die Formel für die Ableitung von  $f(x) = \frac{u(x)}{v(x)}$  lautet nun jedoch  $f'(x) = \frac{u'(x) \cdot v(x) - u(x) \cdot v'(x)}{v^2(x)}$ . Die Funktion  $u(x)$  entspricht dem Zähler, die Funktion  $v(x)$  entspricht dem Nenner.

**Example 72.** Leite die Funktion  $f(x) = \frac{x^3}{e^x}$  ab.

**Example.** Wie suchen uns als erstes die beiden Teilfunktionen  $u$  und  $v$ . Diese sind hier  $u(x) = x^3$  und  $v(x) = e^x$ . Damit können wir nun die Ableitung errechnen.

$$\begin{aligned} f'(x) &= \frac{u'(x) \cdot v(x) - v'(x) \cdot u(x)}{v^2(x)} \\ &= \frac{3x^2 \cdot e^x - x^3 \cdot e^x}{(e^x)^2} \\ &= \frac{e^x \cdot (3x^2 - x^3)}{e^{2x}} \\ &= \frac{x^2(3 - x)}{e^x} \end{aligned}$$

26.6.4. *Ableiten von Verkettungen — die Kettenregel.* Es gibt Funktionen  $f(x) = u(v(x))$ . Es gibt also zwei Funktionen die man ableiten müsste. Die *Kettenregel* besagt, dass die Ableitung von  $f(x) = u(v(x))$  der Term  $f'(x) = u'(v(x)) \cdot v'(x)$ . Man leitet also zuerst die äußere Funktion  $u(x)$  ab, danach schreibt man als Produkt die Ableitung der inneren Funktion  $v(x)$  dahinter.

**Example 73.** Ableitung von  $f(x) = 3(x - 2x^2)^3$

**Example.** Die äußere Funktion ist  $u(x) = 3x^3$ . Die innere Funktion ist  $v(x) = x - 2x^2$ . Die Ableitung lautet also  $f'(x) = 9(x - 2x^2)^2 \cdot (1 - 4x)$ .

**Exercise 1.** Leite  $f(x) = \frac{1}{\sqrt{x^3 - \ln(x)}}$  ab. Lösung:  $f'(x) = -\frac{3x^3 - 1}{2x \cdot \sqrt{(x^3 - \ln(x))^3}}$ .

26.7. **Tangenten.** Es kommt oft in der Analysis vor, dass man Geraden berechnen soll, die den Graphen einer Funktion nur in einer gegebenen Stelle berühren. Diese Art von Geraden nennt man Tangenten.

Die Gleichung zur Berechnung einer Tangente an der Stelle  $x_0$  des Graphen lautet  $t(x) = f'(x_0)(x - x_0) + f(x_0)$ .

26.7.1. *Normale einer Tangenten.* Es gibt Geraden, die auf andere Geraden senkrecht stehen. Diese bezeichnet man als Normalen.

Einer Normale zu einer Tangente schneidet diese in dem Tangentialpunkt und steht senkrecht auf diese. Die Steigung einer Normalen gegenüber der Steigung  $m$  der Tangente ist  $m' = -\frac{1}{m}$ . Die Gleichung zur Berechnung der Normalen einer Tangente lautet also  $n(x) = -\frac{1}{f'(x_0)}(x - x_0) + f(x_0)$ . Es ist zu beachten, dass der  $y$ -Achsenabschnitt der Normalen im Normalfall nicht der gleiche wie bei der Tangente ist.

26.7.2. *Wendetangenten.* Oftmals soll eine Tangente berechnet werden, die den Graphen an der Stelle des Wendepunktes berührt. Um diese aufzustellen berechnet man erst die Wendepunkte eines Graphen (siehe 26.10), anschließend berechnet man mit  $t_w(x) = f'(x_w)(x - x_w) + f(x_w)$  die Wendetangenten.

26.8. **Übersicht über bekannte Ableitungen.** Tab. 1 ist eine kleine Übersicht der bekanntesten Funktionen und deren Ableitungen.

26.9. **Extremstellen.** Extremstellen liegen dort, wo die Funktion  $f(x)$  den minimalsten oder maximalsten Wert annimmt. An diesen Stellen ist die Steigung der Funktion 0, daher kann man die Extremstellen ermitteln, indem man die erste Ableitung gleich 0 setzt und danach diese Punkte mit der zweiten Ableitung überprüft.

Wenn  $f''(x)$  an einer ermittelten Extremstelle  $> 0$  ist, dann liegt dort ein *Tiefpunkt*. Wenn  $f''(x)$  an einer ermittelten Extremstelle  $< 0$  ist, dann liegt dort ein *Hochpunkt*. Wenn  $f''(x)$  an dieser Stelle jedoch 0 ist, existiert hier keine Extremstelle.

Statt der Überprüfung mit der zweiten Ableitung kann man die Art des Extrempunktes auch durch Untersuchung auf VZW<sup>16</sup> der ersten Ableitung ermitteln. Für einen VZW (von links nach rechts gesehen) gilt:

<sup>16</sup> Vorzeichenwechsel



- VZW von + nach -: *Hochpunkt*
- VZW von - nach +: *Tiefpunkt*

**Example 74.** Extremstellen von  $f(x) = (x+1)^2 - 4$

**Example.** Zuerst muss die erste Ableitung berechnet werden:  $f'(x) = 2(x+1)$ . Die Extremstellen berechnen wir nun mit

$$\begin{aligned} f'(x) &= 0 \\ \Leftrightarrow 2x &= -2 \\ \Leftrightarrow x &= -1 \end{aligned}$$

An der Stelle  $x_0 = -1$  könnte also eine Extremstelle liegen. Dies muss noch mittels der zweiten Ableitung überprüft werden.

$f''(-1) = 2$ . Die zweite Ableitung ist also immer positiv. An der Stelle  $x_0 = -1$  liegt also ein Tiefpunkt. Die alternative Überprüfung auf VZW an der Stelle  $-1$  ergibt einen Wechsel von  $-$  nach  $+$ . Die Untersuchung auf VZW ergibt also auch einen Tiefpunkt.

Die Koordinaten des Tiefpunktes sind  $T(-1 \mid f(-1)) = T(-1 \mid 4)$ .

**26.10. Wendestellen.** Die Berechnung von Wendestellen ist ähnlich wie bei Extremstellen (siehe 26.9). Statt der ersten Ableitung muss die zweite Ableitung  $= 0$  gesetzt werden. Die erhaltenen  $x$  Werte müssen nun mit der dritten Ableitung überprüft werden.

Wenn die dritte Ableitung an der Stelle  $x_0$  einen Funktionswert  $\neq 0$  ergibt, so existiert an dieser Stelle  $x_0$  eine Wendestelle.

**Example 75.** Wendestellen der Funktion  $f(x) = 3x^4 - 3x^2 + 4$ .

**Example.** Die zweite Ableitung lautet  $f''(x) = 36x^2 - 6 = 6 \cdot (6x^2 - 1)$ . Die dritte Ableitung lautet  $f'''(x) = 72x$ .

$f''(x) = 0$  ergibt  $x_0 = -\frac{1}{\sqrt{6}}$  und  $x_1 = \frac{1}{\sqrt{6}}$ . Die Überprüfung für  $x_0$  ergibt  $f'''(-\frac{1}{\sqrt{6}}) < 0$  und  $f'''(\frac{1}{\sqrt{6}}) > 0$ .

Bei beiden  $x$  Werten liegen also Wendestellen. Die Wendestellen sind  $W_1(-\frac{1}{\sqrt{6}} \mid 4,35)$  und  $W_2(\frac{1}{\sqrt{6}} \mid 5,36)$

**26.11. Ganzrationale Funktionen.** In ganzrationale Funktionen tauchen keine Brüche auf. So ist zum Beispiel die Funktion  $f(x) = x^3$  eine ganzrationale Funktionen.

Die allgemeine Form der ganzrationalen Funktion lautet  $f(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0$ .

**26.12. Komplette Funktionsuntersuchung ganzrationaler Funktionen.**

**26.12.1. Symmetrie.** Symmetrie ist in Kapitel 26.1 beschrieben.

**26.12.2. Verhalten.** Das Verhalten für  $x \rightarrow \infty$  und für  $x \rightarrow -\infty$  soll untersucht werden. Das Berechnung ist in Kapitel 26.2 erklärt.

**26.12.3. Schnittpunkt mit y-Achse.** Es soll der Funktionswert  $f(0)$  berechnet werden. Dies vereinfacht das Zeichnen des Graphen.

**26.12.4. Nullstellen.** Es sollen die Nullstellen der Funktion berechnet werden. Dies ist in Kapitel 26.5 erklärt.

**26.12.5. Ableitungen.** Es sollen die ersten 3 Ableitungen der Funktion berechnet werden. Das Ableiten von Funktionen ist in Kapitel 26.6 erklärt.

| <i><b>Funktion</b></i> | <i><b>Ableitung</b></i>        |
|------------------------|--------------------------------|
| $f(x) = C$             | $f'(x) = 0$                    |
| $f(x) = mx + b$        | $f'(x) = m$                    |
| $f(x) = x^2$           | $f'(x) = 2x$                   |
| $f(x) = x^3$           | $f'(x) = 3x^2$                 |
| $f(x) = \frac{1}{x}$   | $f'(x) = -\frac{1}{x^2}$       |
| $f(x) = a^x$           | $f'(x) = a^x \cdot \ln(a)$     |
| $f(x) = e^x$           | $f'(x) = e^x$                  |
| $f(x) = \sin(x)$       | $f'(x) = \cos(x)$              |
| $f(x) = \cos(x)$       | $f'(x) = -\sin(x)$             |
| $f(x) = \tan(x)$       | $f'(x) = \frac{1}{\cos^2(x)}$  |
| $f(x) = \cot(x)$       | $f'(x) = -\frac{1}{\sin^2(x)}$ |

TABELLE 1. Bekannte Ableitungen



26.12.6. *Extremstellen.* Es sollen die Extremstellen der Funktion berechnet werden. Wie dies funktioniert ist in Kapitel 26.9 erklärt.

26.12.7. *Wendestellen.* Es sollen die Wendestellen der Funktion berechnet werden. Wie dies funktioniert ist in Kapitel 26.10 erklärt.

26.13. **Gebrochenrationale Funktionen.** Gebrochenrationale Funktionen enthalten Brüche. Deshalb treten in gebrochenrationalen Funktionen einige Besonderheiten auf, die bei ganzrationalen Funktionen nicht zu beobachten waren. Dazu zählen z.B. Polstellen (siehe Kapitel 26.3) oder senkrechte Asymptoten (siehe Kapitel 26.4).

26.14. **Komplette Funktionsuntersuchung gebrochenrationaler Funktionen.** Dieser Abschnitt zeigt, welche Dinge in welcher Reihenfolge berechnet werden müssen, um eine komplette Funktionsuntersuchung einer gebrochenrationalen Funktion durchzuführen.

26.14.1. *Definitionsmenge.* Die Definitionsmenge einer gebrochenrationalen Funktion wird durch die Nullstellen des Nenners bestimmt.

26.14.2. *Symmetrie.* Symmetrie ist in Kapitel 26.1 beschrieben.

26.14.3. *Polstellen, senkrechte Asymptoten.* Polstellen und senkrechte Asymptoten sind in Kapitel 26.3 und Kapitel 26.4 erklärt.

26.14.4. *Verhalten.* Die Berechnung des Verhalten ist in Kapitel 26.2 erklärt.

26.14.5. *Schnittpunkt mit y-Achse.* Es soll der Funktionswert  $f(0)$  berechnet werden. Dies vereinfacht das Zeichnen des Graphen.

26.14.6. *Nullstellen.* Wie in Kapitel 26.5 erklärt.

26.14.7. *Ableitungen.* Bei den Ableitungen muss meistens die Quotientenregel angewandt werden.

26.14.8. *Extremstellen.* Wie in Kapitel 26.9 erklärt.

26.14.9. *Wendestellen.* Wie in Kapitel 26.10 erklärt.

26.15. **Exponentialfunktionen.** Bei Exponentialfunktionen steht das  $x$  in der Hochzahl, also zum Beispiel  $f(x) = 3^x$ . Die Exponentialfunktion  $f(x) = a^x$  leitet man ab zu  $f'(x) = a^x \cdot \ln(a)$ .

**Example 76.** Leite die Funktion  $f(x) = a^{3x^2}$  ab.

**Example.** In diesem Fall muss die Kettenregel (Kapitel 26.6.4) angewandt werden. Die Ableitung lautet also:  $f(x) = a^{3x^2} \cdot \ln(a) \cdot 6x = 6x \cdot \ln(a) \cdot a^{3x^2}$ .

26.15.1. *Die e Funktion.* Ein besonderer Fall der Exponentialfunktion  $a^x$  ist die sogenannte  $e$ -Funktion. Sie hat statt dem  $a$  die Konstante  $e$  in der Gleichung stehen, sie sieht also folgendermaßen aus:  $f(x) = e^x$ . Das besondere an der Konstanten  $e$  ist, dass ihr Logarithmus eins ist:  $\ln(e) = 1$ . Dadurch vereinfachen sich also die Ableitungen der  $e$  Funktion:  $f'(x) = e^x \cdot \ln(e) = e^x$ . Egal wie oft man die Funktion  $e^x$  also differenziert, sie bleibt immer  $e^x$ .

Umgekehrt bleibt das Integral von  $f(x) = e^x$  immer  $F(x) = e^x$ .

26.16. **Logarithmusfunktionen.** Mit einem Logarithmus kann man die Lösung von Gleichungen wie  $a^x = 5$  berechnen. Die Lösung ist in diesem Falle  $x = \ln(5)$ .

Zu den Logarithmusfunktionen lässt sich allgemein sagen, dass sie die Umkehrfunktion der Exponentialfunktionen sind. Ein Logarithmus hat immer eine Basis:  $x =_{\text{basis}} \log(a)$ . Diesen kann man berechnen durch  $x = \frac{\ln(a)}{\ln(\text{basis})}$ .

So kann man zum Beispiel die Gleichung  $15^x = 5$  mit dem Logarithmus zur Basis 15 berechnen:  $x =_{15} \log(5) = \frac{\ln(5)}{\ln(15)} = 0,594$ . Zur Kontrolle können wir nun  $15^{0,594}$  berechnen. Wenn hier 5 herauskommt, ist das Ergebnis richtig.

26.16.1. *Der 2er Logarithmus —  $\lg(x)$ .* Der 2er Logarithmus hat immer die Basis 2. Seine Schreibweise ist  $x = \lg(a)$  Es kann auch geschrieben werden als  ${}_2\log(a) = x$ .

26.16.2. *Der 10er Logarithmus —  $\lg(x)$ .* Der 10er Logarithmus hat immer die Basis 10. Seine Schreibweise ist  $x = \lg(a)$  Es kann auch geschrieben werden als  ${}_{10}\log(a) = x$ .

26.16.3. *Der natürliche Logarithmus —  $\ln(x)$ .*  $\ln(x)$  wird bezeichnet als »logarithmus naturalis«, er hat die Basis  $e$ . Man schreibt ihn als  $x = \ln(a)$ .

26.16.4. *Gesetze des Logarithmus.*

- (1)  $\ln(a \cdot b) = \ln(a) + \ln(b)$
- (2)  $\ln\left(\frac{a}{b}\right) = \ln(a) - \ln(b)$
- (3)  $\ln(a^b) = b \cdot \ln(a)$

## 27. INTEGRALRECHNUNG

**27.1. Allgemeines.** Es gibt zu jeder Funktion  $f(x)$  eine sogenannte Stammfunktion  $F(x)$ , die abgeleitet wieder die ursprüngliche Funktion  $f(x)$  ergibt. Diese Stammfunktionen werden in der Integralrechnung berechnet. Die Stammfunktion berechnet sich mit  $F(x) = \int f(x) dx$ , gesprochen »Groß  $F(x)$  ist gleich Integral von  $f(x) dx$ «.

**27.1.1. Allgemeine Regeln.** Die Stammfunktion einer Potenzfunktion wie  $f(x) = x^2$  ist  $F(x) = \int x^2 dx = \frac{1}{3}x^3$ . Allgemein geschrieben ist die Stammfunktion von  $f(x) = x^n$  die Funktion  $F(x) = \frac{1}{n+1}x^{n+1}$ . Die Potenz wird also um 1 erhöht und anschließend wird der Kehrwert der neuen Potenz vor das Integral geschrieben. Es gilt weiterhin:

- $\int_a^b f(x) dx = -\int_b^a f(x) dx$ . Bei Vertauschen der Integrationsgrenzen ändert sich also das Vorzeichen des Integrals
- $\int_a^b f(x) dx + \int_a^b g(x) dx = \int_a^b (f(x) + g(x)) dx$ . Integrale mit gleichen Integrationsgrenzen kann man also addieren.
- $\int c \cdot f(x) dx = c \cdot \int f(x) dx$ . Konstante Faktoren darf man vor das Integral schreiben, da sie unverändert bleiben.

**27.2. Produktintegration.** Mit der Produktintegration ist es einfacher, komplexe Produkte zu integrieren. Die allgemeine Formel lautet:

$$\int uv' = uv - \int u'v. \text{ Dies lässt sich umformen nach } \int u dv = uv - \int v du.$$

**Example 77.** Produktintegration

**Example.**

$$\int x \cdot \ln(x) dx = \int \ln(x) \cdot x dx = \ln(x) \cdot \frac{x^2}{2} - \int \frac{1}{x} \cdot \frac{x^2}{2} dx = \frac{1}{2}x^2 \cdot \left(\ln(x) + \frac{1}{2}\right) + C$$

Nicht vergessen, die Integrationskonstante  $C$  hinzuzufügen!

**27.3. Tipps und Tricks zum Integrieren.**

- Zur Produktintegration:
  - Wenn ein Logarithmus im Integral steht, sollte dieser normalerweise dort stehen bleiben. Der andere Faktor muss also ins Differential geschrieben werden (d.h. dieser Faktor ist dann  $v'$ ).
  - Wenn die  $e$  Funktion im Integral steht, wird diese normalerweise als  $v'$  angenommen und ins Differential geschrieben.
  - In Fällen wie  $\int \ln(x) dx$  ist es sinnvoll, die Funktion  $u'(x) = 1$  zu setzen. Die Lösung dieses Integrals wäre  $x \cdot (\ln(x) - 1)$ .
- Integrale der Form  $\int e^{k \cdot x} dx$  oder  $\int \cos(k \cdot x) dx$  werden abgeleitet zu  $\frac{1}{k}e^{k \cdot x}$  und  $\frac{1}{k}\sin(k \cdot x)$ . Der Faktor  $k$  wird also als Kehrwert davorgeschrieben.

**27.4. Integration durch Substitution.** Es ist die Umkehrung der Kettenregel als Differentiationsregel.

$$(22) \quad \int g(u) \cdot u'(x) dx = \int g(u) du|_{u=u(x)}$$

Aussprache des rechten Teils: »Integral über  $g$  von  $u$  nach  $du$  für  $u = u(x)$ « (d.h. es ist Rückersetzung nötig).

**Example 78.** Integration durch Substitution von  $\int (3 + x^5) x^4 dx$

**Example.** Zuerst forme man den Integranden so um, dass er aus dem Produkt einer verketteten Funktion  $g(u(x))$  und der Ableitung der inneren Funktion  $u'(x)$  besteht; anschließend wendet man Gleichung 22 an:

$$\begin{aligned}
 \int (3 + x^5) x^4 dx &= \frac{1}{5} \int (3 + x^5) 5x^4 dx \\
 &= \frac{1}{5} \int 3 + u du|_{u=x^5} \\
 &= \frac{1}{5} \frac{(3 + u)^2}{2} + C|_{u=x^5} \\
 &= \frac{(3 + x^5)^2}{10} + C
 \end{aligned}$$

Formale Vorgehensweise bei der Substitution:

- (1) Sei  $\int g(u(x)) \cdot u'(x) dx$
- (2) Setze  $u = u(x)$
- (3)  $\frac{du}{dx} = u'(x)$  woraus  $du = u'(x) dx$  wird
- (4) Dann ersetze  $u(x)$  durch  $u$  und  $u'(x) dx$  ersetzen.

**27.5. Übersicht über bekannte Integrale und Stammfunktionen.** Tab. 2 ist eine kleine Übersicht über die bekanntesten Funktionen und deren Stammfunktionen.

| Funktion                     | Stammfunktion                            |
|------------------------------|------------------------------------------|
| $f(x) = k$                   | $F(x) = k \cdot x + C$                   |
| $f(x) = x^n$                 | $F(x) = \frac{1}{n+1} \cdot x^{n+1} + C$ |
| $f(x) = \frac{1}{x}$         | $F(x) = \ln(x) + C$                      |
| $f(x) = \sin(x)$             | $F(x) = -\cos(x) + C$                    |
| $f(x) = \cos(x)$             | $F(x) = \sin(x) + C$                     |
| $f(x) = \frac{1}{\cos^2(x)}$ | $F(x) = \tan(x) + C$                     |
| $f(x) = e^x$                 | $F(x) = e^x + C$                         |
| $f(x) = a^x$                 | $F(x) = \frac{1}{\ln(a)} \cdot a^x + C$  |
| $f(x) = \sqrt{x}$            | $F(x) = \frac{2}{3} \sqrt{x^3} + C$      |

TABELLE 2. Bekannte Stammfunktionen

**27.6. Flächenberechnung.** Mit Hilfe von Integralen kann man die Fläche berechnen, die ein Graph umschließt.

**27.6.1. Fläche zwischen Graph und  $x$ -Achse.** Die einfachste Art der Flächenberechnung ist die Berechnung der Fläche zwischen Graph und  $x$ -Achse. Der Bereich wird durch die Integrationsgrenzen des Integrals angegeben.

Es müssen zuerst die Nullstellen der Integralfunktion berechnet werden. Wenn nämlich eine Nullstelle innerhalb des Integrationsbereiches liegt, so muss sich die Gesamtfläche aus den Teilstücken von Nullstelle zu Nullstelle zusammensetzen. Die Teilstücke müssen in Absolutbeträgen stehen, da eine Fläche nicht negativ sein kann! Bei Funktionen wie  $A_1 = \int_{-1}^1 x^3 dx$  würde die Fläche bei normaler Integration nämlich 0 betragen! Bei Berechnung mit  $A_2 = \left| \int_{-1}^0 x^3 dx \right| + \left| \int_0^1 x^3 dx \right|$  beträgt sie richtigerweise  $A_2 = 2 \text{ FE}$ .

**Example 79.** Fläche zwischen dem Graph der Funktion  $f(x) = \frac{1}{5}x \cdot e^x$  in den Grenzen  $a = -2$  und  $b = 3$ .

**Example.** Die einzige Nullstelle der Funktion ist  $x_0 = 0$ . Diese liegt im Integrationsbereich, daher müssen zwei Einzelintegrale berechnet werden. Die Stammfunktion ist

$$F(x) = \int \frac{1}{5} x \cdot e^x dx = \frac{1}{5} [e^x(x - 1)]$$

Die eingeschlossene Fläche ist also

$$\begin{aligned}
 A_{ges} &= A_1 + A_2 \\
 &= \frac{1}{5} |e^x(x - 1)|_{-2}^0 + \frac{1}{5} |e^x(x - 1)|_0^3 \\
 &= \frac{1}{5} \left( \left| \frac{3}{e^2} - 1 \right| + |2e^3 + 1| \right) \\
 &= 8,35 \text{ FE}
 \end{aligned}$$

**27.6.2. Fläche zwischen zwei Graphen.** Es kann passieren, dass man die Fläche berechnen soll, die zwischen zwei Graphen der Funktionen  $f(x)$  und  $g(x)$  eingeschlossen werden. Um dies zu berechnen stellt man die Differenzfunktion  $d(x) = f(x) - g(x)$  auf. Mit dieser Funktion  $d(x)$  kann man den Flächeninhalt berechnen, der zwischen den Graphen eingeschlossen wird. Der Flächeninhalt beträgt dann  $A = \left| \int d(x) dx \right|$ . Die Betragsstriche sind notwendig, da das Integral über der Differenzfunktion auch einen negativen Wert annehmen kann. Ein Flächeninhalt kann jedoch nur positiv sein.

Um die Gesamtfläche zu berechnen, muss man von Schnittpunkt zu Schnittpunkt integrieren. Man berechnet also die Schnittpunkte mit Hilfe der Gleichung  $f(x) = g(x)$ , d.h. man berechnet die Nullstellen der Differenzfunktion  $d(x)$ . Diese Punkte geben die Integrationsgrenzen an. Wenn die beiden Funktionen z.B. 3 Schnittpunkte haben, muss zuerst vom ersten zum zweiten und dann vom zweiten zum dritten Schnittpunkt integriert werden. Man darf auch hier nicht vergessen die Absolutbeträge um die Teilintegrale zu setzen.

**Example 80.** Berechnung der Fläche zwischen den Graphen  $f(x) = -x^2 + 5$  und  $g(x) = x^2$

**Example.** Die Differenzfunktion  $d(x)$  lautet  $d(x) = f(x) - g(x) = -x^2 + 5 - x^2$ . Die Schnittpunkte der beiden Graphen sind bei  $x_1 = -\sqrt{2,5}$  und  $x_2 = \sqrt{2,5}$ . Das Flächeninhalt beträgt also

$$A = \left| \int_{-\sqrt{2,5}}^{\sqrt{2,5}} -2x^2 + 5 dx \right| = 10,54 FE$$

**27.7. Volumenberechnung.** Die Flächenberechnung kann zur Volumenberechnung ausgeweitet werden. Im Normalfall soll das Volumen eines Körpers berechnet werden, der entsteht, wenn der Graph einer Funktion um die  $x$  oder  $y$  Achse rotiert wird. Die benötigte Formel zur Volumenberechnung bei Rotation um die  $x$ -Achse lautet

$$V = \pi \cdot \int_a^b f^2(x) dx$$

Der Körper hat die Grenzen  $a$  und  $b$ . Es darf wie bei der Flächenberechnung nicht vergessen werden, von Nullstelle zu Nullstelle zu integrieren! Auch darf man nicht vergessen, um die Teilintegrale Betragsstriche zu setzen!

Wenn der Körper um die  $y$ -Achse rotiert werden soll, dann gilt folgende Gleichung

$$V = \pi \cdot \int_c^d \bar{f}^2(x) dx = V = \pi \cdot \int_c^d x^2 dy$$

Es wird also einfach die Umkehrfunktion der Funktion benutzt.

Wenn ein Körper rotiert werden soll, der zwischen zwei Funktionen  $f(x)$  und  $g(x)$  entsteht, dann werden die Funktionen erst quadriert und ihre Differenzfunktion anschließend rotiert:

$$V = \pi \cdot \left| \int g^2(x) - f^2(x) dx \right|$$

Es darf hier wieder nicht vergessen werden, die Betragsstriche zu setzen und von Schnittpunkt zu Schnittpunkt zu integrieren!

**27.8. Länge eines Bogens.** Die Länge eines Bogens gibt die Länge an, die eine Funktion als »Oberfläche« hat. Die Formel zur Berechnung der Länge lautet  $l = \int_a^b \sqrt{1 + [f'(x)]^2} dx$ . Wenn die Ableitung  $f'(x)$  die Variable  $x$  enthält, wird das Integral sehr kompliziert. Zur Berechnung dient folgendes Integral aus der Formelsammlung, das hier nicht bewiesen werden kann:

$$\int \sqrt{x^2 + a^2} dx = \frac{x}{2} \sqrt{x^2 + a^2} + \frac{a^2}{2} \cdot \ln \left| \frac{x}{a} + \frac{1}{a} \sqrt{x^2 + a^2} \right|$$

Die Variable  $x$  erhält natürlich denjenigen Teil des Integrales, der auch die  $x$ -Variable enthält.

## 28. PROF. LÖFFLER: ABBILDUNGEN VON MENGEN («KOMMUNIKATION VON MENGEN»)

Funktionen bestehen aus einer Rechenvorschrift (die rechte Seite), durch die für jeden  $x$ -Wert ein  $y$ -Wert berechnet wird (so dass sich die Möglichkeit einer Auftragung mit den Wertpaaren  $(x, y)$  ergibt. Dies ist eine Abbildung  $x \mapsto y$ , wobei  $x \neq -3$  sein muss.

Die Verwendung von Zahlen für Abbildungen hat den Vorteil, dass man mit Zahlen rechnen kann. Eine Abbildung nämlich (nur) bei Zahlenmengen als Rechenvorschrift angegeben werden. Man könnte auch z.B. Farben als Elemente für Abbildungen verwenden, muss dann aber die Abbildung einzeln für jedes Element angeben (Pfeile im Mengendiagramm). Mit den Farben selbst jedoch lässt sich nicht rechnen. Durch Farbmodelle (CMYK, RGB) wird einem Zahlen- $n$ -Tupel eine Farbe zugeordnet, jedoch wird nicht unbedingt jede Farbe des Farbraums mit einem solchen Zahlen- $n$ -Tupel versehen.

Eine vollständige Abbildung muss immer über drei Dinge definiert werden; nur wenn bei zwei Abbildungen alle drei Teile gleich sind, sind diese Abbildungen identisch.

- (1) Die Grundmenge
- (2) Die Bildmenge
- (3) die Abbildungsvorschrift

Deshalb beginnt eine Kurvendiskussion mit der Bestimmung von Definitions- und Wertemenge, um die Rechenvorschrift (die Funktion) als Abbildung auffassen zu können. Beispiel für eine vollständige Abbildung:

$$\begin{aligned} \mathbb{N} &\longrightarrow \mathbb{N} \cup \{-1\} \\ s(n) &= n - 1 \\ n &\mapsto n - 1 \end{aligned}$$

Abbildungen sind in der Informatik wichtig, weil sie eine universelle Sprache zur Beschreibung von Beziehungen darstellt (Komposition usw.).

### Definition 115. Abbildungen

**Definition.** Eine Abbildung bedeutet, dass jedem Element aus  $M$  ein Element aus  $N$  zugeordnet wird.

Funktionen sind solche Abbildungen. Damit wirklich jedem Element aus  $M$  ein Element aus  $N$  zugeordnet wird, muss  $M \subseteq D_{Abb}$  sein ( $D_{Abb}$  ist die Definitionsmenge der Abbildung, die Elemente, für die die Abbildung definiert ist); d.h. die Abbildung muss auf jedes Element von  $M$  anwendbar sein.

### Example 81. Abbildung durch eine Rechenvorschrift

**Example.** Es sei  $\mathbb{N} = \{0, 1, 2, 3, \dots\}$ . Eine Abbildung wird durch einen Pfeil dargestellt:  $a : \mathbb{N} \longrightarrow \mathbb{N}$ . Für einzelne Elemente wird ein Pfeil der Form  $3 \mapsto 4$  verwendet. Eine Abbildung muss allgemein durch eine Abbildungsvorschrift angegeben werden (z.B.  $a(n) = n + 1$ ), die Angabe von Beispielen lässt nur Vermutungen über die Abbildung zu. Wenn die Grundmenge keine Rechenoperatoren beinhaltet, so ist nur eine Angabe über Beispiele / explizite erschöpfend aufzählende Aussagen für jedes Element der Grundmenge möglich.

### Example 82. Abbildung durch eine Rechenvorschrift

**Example.** Es sei die Abbildung  $m : \mathbb{N} \longrightarrow \mathbb{N}$  mit der Abbildungsvorschrift  $m(n) = 2 \cdot n$ .

### Example 83. Die identische Abbildung.

**Example.** Diese Abbildung existiert für jede Menge. Angegeben als:

$$\begin{aligned} M &\longrightarrow M \\ m &\mapsto m \end{aligned}$$

Dazu gehört auch als allgemeinerer Fall bei zwei Mengen  $A, M$  die immer existierende Inklusionsabbildung  $i$  für  $A \subset M$ :

$$\begin{aligned} i : A &\longrightarrow M \\ a &\mapsto a \end{aligned}$$

### Example 84. Die kollabierende Abbildung.

**Example.** Jedes  $m \in M$  wird auf ein Element  $x \in M$  abgebildet.

$$\begin{aligned} M &\longrightarrow M \\ m &\mapsto x \end{aligned}$$

An dieser Abbildung zeigt sich, dass eine Abbildung eine eindeutige, aber nicht eine eineindeutige Zuordnung bedeutet. Weiteres Beispiel:  $\mathbb{Z} \longrightarrow \mathbb{Z}, x \mapsto x^2$ . Aufgrund der eindeutigen Zuordnung darf ein Element aus der Definitionsmenge nicht zwei Elemente aus einer Wertemenge zugeordnet werden. Die Vorschrift »Ordne einer Zahl  $x$  aus  $\mathbb{Z}$  die Menge aller Zahlen zu, deren Quadrat  $x$  ist« beschreibt also keine Abbildung, weil  $1 \mapsto \{-1, 1\}$  möglich wäre.

### 28.1. Komposition (Zusammensetzung) von Abbildungen. Die Abbildung

$$\begin{aligned} M &\xrightarrow{f} N \xrightarrow{g} O \\ m &\mapsto f(m) \mapsto g(f(m)) \end{aligned}$$

kann dargestellt werden durch eine Abbildung  $h(m) = g(f(m))$ . Dies entspricht in der Ausdrucksweise für Funktionen der Kettenregel. So zum Beispiel ist die Verkettung der Abbildungen  $g(x) = x^2$ ,  $h(x) = 2 \cdot x$  die Funktion  $i(x) = g(h(x)) = (2 \cdot x)^2 = 4 \cdot x^2$ . Für die Verkettung von Abbildungen  $g, f$  wird das Zeichen  $\circ$  eingeführt:  $g \circ f$  ist » $g$  kombiniert mit  $f$ «.

### Example 85. Komposition von Abbildungen: Die ASCII-Ordnung

**Example.**

$$UPPER \longrightarrow \{1, 2, 3, \dots, 26\} \longrightarrow \{0, 1, 2, \dots, 255\}$$

Die erste Abbildung sei dabei  $A \mapsto 1$ ,  $B \mapsto 2$  usw., die zweite Abbildung sei dabei  $z \mapsto z + 64$ . Die Verkettung beider Abbildungen ist die sog. ASCII-Ordnung.

**Example 86.** zyklisches Vertauschen

**Example.** Es bestehe die Abbildung

$$\begin{aligned} \{1, 2, 3, \dots, 26\} &\xrightarrow{1} \{1, \dots, 26\} \\ n &\mapsto n + 1 \\ 26 &\mapsto 1 \end{aligned}$$

Diese Abbildung ist das sog. zyklische Vertauschen, darstellbar an einem Kreis mit 26 nummerierten Segmenten, bei dem jedem Segment das folgende zugeordnet wird, also dem Segment 26 die 1. Durch Komposition mit der Abbildung

$$UPPER \longrightarrow \{1, 2, 3, \dots, 26\} \longrightarrow \{0, 1, 2, \dots, 255\}$$

(wobei  $A \mapsto 1$ ,  $B \mapsto 2$  usw.) und Rückabbildung auf die Buchstaben ergibt sich eine einfache Form der Kodierung, bei dem jedem Buchstaben der Folgebuchstabe zugeordnet wird (wurde von Cäsar angewandt; siehe Zeichnung 1). Durch Verschiebung anhand der einzelnen Buchstaben eines stets wiederholten Codewortes kann diese Form der Kodierung verbessert werden.

Arten von Abbildungen:

**injektiv::** eine eindeutige Abbildung. Nicht z.B.:  $q : \mathbb{Z} \longrightarrow \mathbb{Z}$  mit  $q(x) = x^2$ . Die identische Abbildung ist dagegen injektiv. Allgemein formuliert:  $f : M \longrightarrow N$  heißt injektiv, falls  $f(m) = f(n) \Rightarrow m = n$ . Beweis, dass die Abbildung

$$\begin{aligned} \mathbb{Z} &\longrightarrow \mathbb{Z} \\ n &\mapsto 2n \\ f(n) &= 2n \end{aligned}$$

injektiv ist: sei  $f(n) = f(m) \Rightarrow 2n = 2m \Rightarrow n = m$  qed. Diese Abbildung ist jedoch nicht surjektiv, weil nur die geraden ganzen Zahlen getroffen werden.

**surjektiv::**  $f : M \longrightarrow N$  heißt surjektiv, falls zu  $n \in N$  gibt es  $m$  mit  $f(m) = n$ . D.h. jedes Element in der Wertemenge wird tatsächlich von der Abbildung getroffen. Beispiel:  $\mathbb{Z} \longrightarrow \{1\}$ ,  $n \mapsto n$ . Um zu zeigen, dass eine Abbildung nicht surjektiv ist, muss man ein Element finden, das in der Bildmenge enthalten ist, aber nicht von der Abbildung getroffen wird.

**bijektiv::** eine Abbildung, die sowohl injektiv als auch surjektiv ist. Eine Abbildung muss bijektiv sein, um umkehrbar zu sein: an jedem Element der Bildmenge (surjektiv) kommt genau ein Pfeil (injektiv) an. Die Umkehrabbildung wird mit mit

$$f^{-1}$$

(Index über das Zeichen setzen!) bezeichnet. Es gilt:

$$\begin{aligned} M &\xrightarrow{f} N & f^{-1} &\longrightarrow N \\ m &\xrightarrow{f} f(m) & f^{-1} &\mapsto m \end{aligned}$$

**Example 87.**  $\sin(x)$

**Example.** Ist die Abbildung  $\sin : \mathbb{R} \longrightarrow \mathbb{R}$  surjektiv? Nein, denn man kann ein Element der Bildmenge finden, das nicht von der Abbildung getroffen wird (Parallele zur x-Achse, die den Graphen nicht schneidet).

Ist die Abbildung  $\sin : \mathbb{R} \longrightarrow [-1, +1]$  surjektiv? Ja, denn jede Parallele zur x-Achse im Wertebereich schneidet den Graphen, es wird also jedes Element der Bildmenge von der Abbildung getroffen. Jedoch ist diese Abbildung nicht injektiv, da z.B. den x-Werten  $2\pi$  und  $4\pi$  derselbe Wert zugeordnet wird.

Ist die Abbildung  $\sin : [-\frac{\pi}{2}, \frac{\pi}{2}] \longrightarrow [-1, +1]$  nun bijektiv? Ja, denn keinen zwei x-Werten wird derselbe y-Wert zugeordnet; jede Parallele zur x-Achse schneidet den Graphen in Werte- und Bildmenge nur einmal. Also ist diese Abbildung bijektiv, d.h. umkehrbar. Diese Umkehrung der Abbildung  $\sin(x)$  ist  $\arcsin(x)$ , die jeder Zahl zwischen  $-1$ ,  $1$  ein Bogenmaß zuordnet. Ähnlich verhält es sich mit der Einschränkung des Definitionsbereichs, um  $\tan(x)$  umkehrbar zu machen.

## 29. PROF. LÖFFLER: ZAHLEN

## 29.1. Zahlenräume.

$\mathbb{N} = \{0, 1, 2, \dots\}$ : Die natürlichen Zahlen. Auf ihrer Existenz baut die moderne Mathematik überhaupt auf; deshalb wird ihre Existenz nicht in Frage gestellt. Mögliche Rechenoperationen: Addition, Multiplikation, Division mit Rest ( $\frac{a}{b} = q \text{ Rest } r \Leftrightarrow a = b \cdot q + r$  mit  $0 \leq r < b$ ). Division mit Rest geschieht durch Zählen, wie oft man den Divisor vom Dividenten abziehen kann. Die Subtraktion ( $a - b$ ) ist nur möglich, wenn  $b \leq a$ . Der hier erforderliche Übergang zu negativen Zahlen wurde lange nicht verstanden, denn man konnte sich keine negativen Mengen vorstellen (Mengen »kleiner als nichts«); diese traten beim Zählen nämlich nicht auf, kommen jedoch beim Messen (z.B. von Temperaturen) vor.

$\mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\}$ : Die Erweiterung von  $\mathbb{N}$  um den negativen Bereich.

**29.2. Darstellung von Zahlen.** Beim Zählen konnte man sich durch bijektive Abbildungen zur Darstellung großer Zahlen behelfen: für jedes Schaf einen Kieselstein zur Repräsentation. Dann wurden Zahlen als Strichmengen aufgeschrieben und wegen der Unübersichtlichkeit dieser Darstellung durch neue Zeichen abgekürzt, z.B. *IIIII* durch *V* bei den Römern. Das Römische Zahlensystem basiert darauf: *MMCMXLXIV* usw. Die Addition von derart dargestellten Zahlen war schon enorm aufwändig.

In Mesopotamien wurden Positionen zur Zahldarstellung eingeführt, d.h. eine Vorform des Stellenwertsystems. Grad werden deshalb heute noch mit Positionskennzeichnungen dargestellt:  $93^\circ 21' 12''$ . Nicht besetzte Positionen wurden nicht geschrieben, denn man kannte keine Null. Die Araber erfanden im Mittelalter diese Null, so dass unser Stellenwertsystem entstand:

$$347 = 3 \cdot 10^2 + 4 \cdot 10^1 + 7 \cdot 10^0$$

Nicht besetzte Positionen wurden durch 0 dargestellt; die Leistung war hier, »nichts« trotzdem darzustellen. Dieses positionale Zahlensystem vereinfachte das Rechnen enorm und setzte sich in der breiten Bevölkerung durch; es wurde um 1400 entwickelt. 10 wurde als Basis genommen, weil der Mensch eben einfach 10 Finger hat. Unsere Zahlensymbole  $1, \dots, 9, 0$  sind eine Abkürzung für die Aufzählung mit einer entsprechenden Menge Strichen.

Abstrakte Beschreibung: Für die  $0, \dots, b-1$  ( $b = 10$ ) werden Symbole eingeführt. Die Zahl 10 wird dann dargestellt als:  $1 \cdot b^1 + 0 \cdot b^0 = 10$ , eine Zusammensetzung von Symbolen im Stellenwertsystem. In jedem Zahlensystem braucht man  $b-1$  Symbole (oder unterscheidbare Zustände), wenn  $b = \text{Basis}$ . Mit jeder Basis  $b$  und einer entsprechenden Menge  $b-1$  von Symbolen kann man Zahlen darstellen; Zahlen lassen sich also stets zu jeder beliebigen Basis darstellen.

Statt 10 als Basis hätte man auch eine beliebige andere Basis wählen können, z.B. 12:  $x_1 \in \{0, 1, \dots, 9, A, B\}$ . Darstellungen von Zahlen in diesem System sind z.B.  $B1A$ ,  $111 = 1 \cdot 12^2 + 1 \cdot 12^1 + 1 \cdot 12^0 = 157$ . Weil gleiche Darstellung wie hier verschiedene Zahlen meinen kann, gilt als Vereinbarung: eine Zahl wird immer mit der Basis 10 angegeben, es sei denn ein Index gibt eine andere Basis an wie z.B.:  $347_{12} = 3 \cdot 12^2 + 4 \cdot 12^1 + 7 \cdot 12^0 = 487_{10} = 487$ . Die Basis selbst wird immer in einem Zahlensystem mit Basis 10 angegeben.

Wozu braucht man andere Basen für Zahlen?

- Basis 2 in der Informatik. Symbole: 0, 1. Es wären auch beliebige andere zwei unterscheidbare Zustände möglich. Der Vorteil der Basis 2 ist, dass etwas braucht, das sich (nur) in zwei Zuständen befinden kann (d.h. ein Bit); dies ist die Basis der Digitaltechnik.  
Ein Datentyp ist eine Vorschrift, wie man eine Information übersetzt in eine Bitfolge. Wie eine Bitfolge interpretiert werden soll, ist prinzipiell beliebig, muss also durch eine Festlegung von außen festgelegt werden, nämlich in der Deklaration der Variablen, aus der das Programm erkennt, wie Bitfolgen zu interpretieren sind. Beispiele:
  - Beim Datentyp BOOLEAN sei als Vorschrift zur Übersetzung in eine Bitfolge *false* = 0 und *true* = 1.
  - INT (Integer, übersetzt in eine Folge von 32 bit) als international vereinbarter Datentyp
  - UNSIGNED CHAR (Zeichen, übersetzt in eine Folge von 8 bit) als international vereinbarter Datentyp. Die Bitfolgen werden interpretiert als Binärzahlen, wie Dezimalzahlen von rechts nach links geschrieben. So ist z.B.  $1111111_2 = 255$ , die größte Zahl, die durch 8 bit dargestellt werden kann. Durch eine Abbildung injektive  $ANSI \rightarrow CHAR$  wird jedem Zeichen der ASCII-Tabelle eine Bitfolge CHAR (äquivalent einer Zahl  $\{1, \dots, 127\}$  bei UNSIGNED CHAR) zugeordnet.
  - Eine BITMAP: ein Bild, reihenweise übersetzt in eine Bitfolge, je nachdem, ob ein Bildpunkt schwarz oder weiß sein soll.
- Basis 8 (Oktalsystem) in der Informatik. Symbole:  $0, 1, \dots, 7$ .
- Basis 10. Für Basen  $2, \dots, 10$  sind genügende Symbole vorhanden.
- Basis 16 (Hexadezimalsystem) in der Informatik. Symbole:  $0, 1, 2, \dots, 9, A, B, C, D, E, F$ .
- Basis 12 (Duodezimalsystem). Symbole:  $0, 1, \dots, 9, A, B$

- Basis 26: Symbole:  $A, B, \dots, Z$

**Example 88.** Zahldarstellung in anderen Zahlensystemen

**Example.**

$$1011_2 = 1 \cdot 2^0 + 1 \cdot 2^1 + 0 \cdot 2^2 + 1 \cdot 2^3 = 11$$

$$1011_8 = 1 \cdot 8^0 + 1 \cdot 8^1 + 0 \cdot 8^2 + 1 \cdot 8^3 = 521$$

$$HALLO_{26} = O \cdot 26^0 + L \cdot 26^1 + L \cdot 26^2 + A \cdot 26^3 + H \cdot 26^4$$

$$FF_{16} = F \cdot 16^0 + F \cdot 16^1 = 15 \cdot 16^0 + 15 \cdot 16^1 = 255$$

**Example 89.** Division in unterschiedlichen Zahlensystemen

**Example.** Gegeben ist die Rechnung

$$231_{10} : 10_{10} = 23 \text{ Rest } 1$$

Verwendet man stattdessen die Zahlendarstellung im Neunersystem, ändert sich natürlich das Ergebnis nicht, denn es werden ja dieselben Zahlen verwendet, nur in unterschiedlicher Darstellung:

$$256_9 : 11_9 = 25_9 \text{ Rest } 1$$

### 29.3. Umrechnung zwischen beliebigen Zahlensystemen.

- (1) Gegeben ist eine Zahl in einem beliebigen Zahlensystem.

$$n_b = 231_{10}$$

- (2) Allgemein ist die Darstellung einer Zahl  $n$  zur Basis  $b \geq 2$  in jedem Zahlensystem eindeutig definiert als:

$$(23) \quad n = b_r \dots b_2 b_1 b_0 = b_0 \cdot b^0 + b_1 \cdot b^1 + \dots + b_r \cdot b^r$$

Sei  $b$  nun die Basis des Zielzahlensystems. Der Rest bei ganzzahliger Division  $n \text{ DIV } b$  hat dann die letzte Ziffer der Zahl im Zielzahlensystem als Ergebnis, denn in der Definition der Darstellung einer Zahl (Gleichung 23) ist jeder Summand ohne Rest durch  $b$  teilbar außer  $b_0 \cdot b^0 = b_0 \cdot 1$ , wo der Rest  $b_0$  ist, denn  $b_0 < b$  in einem Stellenwertsystem. Der Rest  $b_0$  nun ist aber die letzte Ziffer im Zielzahlensystem (siehe Gleichung 23). Am Beispiel:

$$231_{10} \text{ DIV } 9 = 23 \text{ Rest } 6$$

6 ist also die letzte Ziffer bei Darstellung von  $231_{10}$  im Neunersystem.

- (3) Der ganzzahlige Anteil des Ergebnisses von  $n \text{ DIV } b$  ist nun die Darstellung der Zahl im Ausgangszahlensystem, jedoch ohne die letzte Ziffer. Also kann man wiederum den Schritt 2 auf dieses Ergebnis anwenden und erhält dann die vorletzte Ziffer im Zielzahlensystem. Dieses Verfahren verwendet man sukzessiv weiter bis die ganze Zahl im Zielzahlensystem vorliegt. Am Beispiel:

$$23_{10} \text{ DIV } 9 = 2 \text{ Rest } 5$$

$$2_{10} \text{ DIV } 9 = 0 \text{ Rest } 2$$

Also ist die Darstellung  $231_{10} = 256_9$ . Der Rest bei Ganzzahldivision  $a \text{ DIV } b$  ist  $a \% b$ , d.h.  $a$  modulo  $b$ . Die Notation mit Prozentzeichen wird in allen gängigen Programmiersprachen verwendet (außer Pascal: MOD). Realisierung im Rechnersystem: nach  $a \text{ DIV } b$  steht das Ergebnis in Register  $a$ , der Rest in Register  $b$ .

Dieses hier beschriebene Verfahren der Konvertierung wird z.B. angewandt bei den Rechner Routinen `atoi`, `itoa` (Konversion ASCII / Integer) in C.

**Example 90.** Anwendung der Umwandlung zwischen Zahlensystemen zur Optimierung von Algorithmen

**Example.** Auszurechnen ist  $2^{100}$  mit möglichst wenig Multiplikationen. Statt 100 Multiplikationen sind bei

$$2^{100} = 2^{64} \cdot 2^{32} \cdot 2^4$$

nur 8 Multiplikationen nötig, nämlich 6 zur Berechnung von  $2^{64}$  und je eine für die letzten beiden Faktoren. Durch Binärdarstellung der Exponenten ergibt sich:

$$2^{1100100_2} = 2^{1000000_2} \cdot 2^{100000_2} \cdot 2^{100_2}$$

Der Vorteil im Binärsystem ist, dass beim Quadrieren einer Zahl nur eine weitere 0 angehängt wird. Also ist die Berechnung der drei obigen Faktoren sehr einfach möglich.



**29.4. Reelle Zahlen: Zahlen zum Messen.** Was zum Beispiel sind die Nachkommastellen der Zahl  $\pi = 3,141592653\dots$ , die zum Beispiel beim Messen des Kreisumfangs  $U = 2\pi r$  vorkommt? Im alten Griechenland waren nur ganze Zahlen bekannt bzw. rationale Zahlen als Brüche ganzer Zahlen (Darstellung von  $\pi$  z.B. als  $\pi = \frac{22}{7}$  oder  $\pi = \frac{355}{113} = 3,141592\dots$ ), jedoch keine Dezimalbrüche zum Messen. Man erkannte durch die Annäherungen der Brüche an  $\pi$  in der Renaissance, dass immer mehr Stellen von Dezimalbrüchen festliegen; so entstand das Konzept der Dezimalbrüche und überhaupt eine Vorstellung von Zahlen:

Zahlen (die Zahlenmenge  $\mathbb{R}$ ) setzen sich zusammen aus allen ganzen Zahlen  $\mathbb{Z}$ , den rationalen Zahlen  $\mathbb{Q}$  (das sind die Brüche  $\frac{p}{q}$  mit  $p, q \in \mathbb{N}$ ) und den irrationalen Zahlen, d.h. den nicht endenden nichtperiodischen Dezimalbrüchen:  $v\dots v, n_1 n_2 n_3 \dots = v\dots v + \frac{n_1}{10^1} + \frac{n_2}{10^2} + \frac{n_3}{10^3} + \dots$  mit  $0 \leq n_i < 1$ . Eine Zahl zu kennen heißt, alle Nachkommastellen zu kennen; das heißt nur, dass man zu einer Zahl (z.B.  $\sqrt{2}$ ) ein Verfahren angeben können muss, um jede beliebige Nachkommastelle dieser Zahl zu berechnen, egal ob dieses Verfahren praktisch durchführbar ist oder nicht. Also kennt man auch  $\pi$ .

Wie kann man nun reelle Zahlen in andere Zahlensystem umrechnen? Die Umrechnung  $x_b \rightarrow x$  geschieht wie bei ganzen Zahlen:

$$x_b = 0, b_1 b_2 b_3 \dots_2 = (b_1 \cdot 10^{-1} + b_2 \cdot 10^{-2} + b_3 \cdot 10^{-3} + \dots)_{10}$$

$$(0,101001)_2 = 1 \cdot \frac{1}{2} + 0 \cdot \frac{1}{4} + 1 \cdot \frac{1}{8} + 0 \cdot \frac{1}{16} + 0 \cdot \frac{1}{32} + 1 \cdot \frac{1}{64} = \frac{32 + 8 + 1}{64} = \frac{41}{64}$$

Umrechnung von Dezimalbrüchen zur Basis 10 in eine andere Basis: Voraussetzung ist zur Umrechnung  $x_{10} \rightarrow x_2$ , dass jede Zahl  $x_{10}$  darstellbar ist als:  $x_{10} = \frac{b_1}{2} + \frac{b_2}{2^2} + \dots$   $b_1$  ist entweder 1 oder 0; es ist 1, wenn die Zahl  $x_{10} \geq \frac{1}{2}$  ist, also wenn  $2x \geq 1$ .

Verfahren in beschreibender Sprache:

- (1) Gegeben:  $0 \leq x < 1$
- (2) Gesucht: Binärdarstellung  $(0, b_1 b_2 b_3 \dots)_2$
- (3) Vorgehensweise: Bilde  $y = 2x$ 
  - (a) Fall  $y \geq 1$ : es ist  $b_1 = 1$ ; setze  $x = y - 1$  und wiederhole Schritt 3
  - (b) Fall  $y < 1$ : es ist  $b_1 = 0$ ; setze  $x = y$  und wiederhole Schritt 3

Verfahren im Pseudocode:

```

Input: x, 0 ≤ x < 1
Output: b1b2b3... eine Folge bi ∈ {0,1}
Variable: i, y
setze i=1
label:
    setze y=2x
    ist y ≥ 1
        setze bi=1
        setze x=y-1
    andernfalls
        setze bi=0
        setze x=y
    setze i=i+1
    ist x=0 breche ab
    gehe zu label

```

Verfahren in C:

```

// keine Konstante, sondern vom Compiler ersetzter Text:
#define MAXLAENGE 50
void binaerdarstellung(double x)
{
    int i=1;
    int double y;
    // der Algorithmus soll spätestens abbrechen, wenn MAXLAENGE
    // Dezimalstellen von x berechnet wurden
    while (x>0.0 && i < MAXLAENGE)
    {
        y=2x;
        if (y>=1.0)
        {
            putchar('1');
            x=y-1.0;
        }
    }
}

```

```

    else
    {
        putchar('0');
        x=y;
    }
    i++;
}
putchar('\n');
}

```

Anwendung dieses Verfahrens zur Umwandlung von Dezimalbrüchen in beliebige Zahlensysteme  $x_{10} \rightarrow x_b$ : Die einzige Änderung ist, dass  $y = bx$  gesetzt wird und als  $b_1$  die entstehende ganze Zahl vor dem Koma verwendet wird.

**Example 91.** Umrechnung eines Dezimalbruchs in Binärdarstellung

**Example.**  $x = 0,625$

$$y = 2x = 1,25 \geq 1$$

Da  $y \geq 1$ , ist  $b_1 = 1$ ; setze  $x = y - 1 = 0,25$

$$y = 2x = 0,5 < 1$$

Da  $y < 1$ , ist  $b_2 = 0$ ; setze  $x = y = 0,5$

$$y = 2x = 1 \geq 1$$

Da  $y \geq 1$ , ist  $b_3 = 1$ ; setze  $x = y - 1 = 0$

Abbruch des Verfahrens mit dem Ergebnis  $x = 0,625 = 0,101_2$

Warum

Durch Verwendung von Brüchen (z.B.  $x_{10} = \frac{1}{5}$ ) können periodische Binärdarstellungen entstehen, wenn nämlich irgendwann der Fall auftritt, dass man im Algorithmus wieder  $x_{10} = \frac{1}{5}$  wird und also die Periode wieder von vorn beginnt.

**Example 92.** Umwandlung von  $\frac{1}{7}$  zur Basis 3

**Example.**  $x = \frac{1}{7}$

$$y = 3x = \frac{3}{7} < 1$$

Da  $y < 1$ , ist  $b_1 = 0$ ; setze  $x = y = \frac{3}{7}$

$$y = 3x = \frac{9}{7}$$

Da  $y \geq 1$ , ist  $b_2 = 1$ ; setze  $x = y - 1 = \frac{2}{7}$

$$y = 3x = \frac{6}{7}$$

Da  $y < 1$ , ist  $b_3 = 0$ ; setze  $x = y = \frac{6}{7}$

$$y = 3x = \frac{18}{7} = 2\frac{4}{7}$$

Da  $y \geq 1$ , ist  $b_4 = 2$ ; setze  $x = y - 1 = \frac{12}{7}$

$$y = \frac{12}{7} = 1\frac{5}{7}$$

Da  $y \geq 1$ , ist  $b_5 = 1$ ; setze  $x = y - 1 = \frac{5}{7}$

$$y = 3x = \frac{15}{7} = 2\frac{1}{7}$$

Da  $y \geq 1$ , setze  $b_6 = 2$ ;

Jetzt wiederholt sich die Periode, da der Rest gleich dem Ausgangsbruch  $\frac{1}{7}$  ist. Also ist die Darstellung:  
 $x = \frac{1}{7} = 0,010212_3$

Die Multiplikation mit 2 im Binärsystem ist sehr einfach: das Bitmuster wird um eine Stelle nach links verschoben (in C:  $n = n << 1$ , der Shift-left-Operator). Ebenso kann die Division durch 2 durch Verschieben des Bitmusters um eine Stelle nach rechts realisiert werden;  $m = m >> 2$  mit dem Shift-right-Operator ist also die Division durch 4. Man sollte diese Notationen um der Effizienz willen verwenden.

Beweis: Sei

$$x = (0, b_1 b_2 b_3 \dots)_2 = b_1 \cdot 10^{-1} + b_2 \cdot 10^{-2} + b_3 \cdot 10^{-3} + \dots$$

so ist

$$2x = (b_1, b_2 b_3 \dots)_2 = b_1 \cdot 10^0 + b_2 \cdot 10^{-1} + b_3 \cdot 10^{-2} + \dots$$

**29.5. Anwendung der Binärdarstellung.** Gleichungen der Form  $10^x = 2$  lassen sich mit den vier Grundrechenarten und Quadratwurzeln lösen; dies entspricht der Berechnung von Logarithmen  $\log 2$ . Wie geht das? Folgende Ausdrücke sind mit den vorausgesetzten Operationen berechenbar:

$$a^{0,12} = a^{\frac{1}{2}} = \sqrt{a}$$

$$a^{0,012} = a^{\frac{1}{4}} = \sqrt{\sqrt{a}}$$

Es ist

$$a^{(0,b_1)_2} = a^{(0,1)_2 \cdot b_1} = \left(a^{(0,1)_2}\right)^{b_1}$$

Denn: die Multiplikation mit  $b_1$  bewirkt

- für  $b_1 = 0$ , dass  $a^{(0,1)_2 \cdot b_1} = a^{(0,1)_2 \cdot 0} = a^0 = a^{(0,b_1)_2} = 1$  wird
- für  $b_1 = 1$ , dass  $a^{(0,1)_2 \cdot b_1} = a^{(0,1)_2 \cdot 1} = a^{(0,1)_2} = a^{(0,b_1)_2} = 1$  wird

Jeder Binärbruch (bzw. allgemein ein Dezimalbruch) lässt sich darstellen als:

$$0, b_1 b_2 \dots b_n = 0, b_1 + 0, 0b_2 + \dots 0, 0 \dots 0b_n$$

Für Potenzausdrücke ergibt sich damit:

$$\begin{aligned} a^{(0,b_1 b_2 \dots b_n)_2} &= a^{0, b_1 + 0, 0b_2 + \dots 0, 0 \dots b_n} = a^{0, b_1} \cdot a^{0, 0b_2} \cdot a^{0, 00b_3} \cdot \dots \cdot a^{0, 0 \dots 0b_n} \\ &= (a^{0,1})^{b_1} + (a^{0,01})^{b_2} + \dots + (a^{0,0 \dots 1})^{b_n} \end{aligned}$$

Somit lässt sich jeder Ausdruck  $a^b$  berechnen über die Umformung des Exponenten  $b$  in Binärdarstellung  $(b)_2$ . Bricht die Binärdarstellung nicht ab, so endet der Algorithmus nicht. Man bricht dann an einer bestimmten Stelle mit einem Resultat ab und weiß: *Resultat*  $\leq$  *Wahrheit*  $\leq$  *Resultat*  $\cdot$  *aktueller Registerwert*. Denn: der wahre Wert entstünde durch Addition weiterer Teile  $(a^{0,0 \dots 1})^{b_n}$  mit den verbleibenden Dezimalstellen, also ist das derzeitige Resultat kleiner. Der aktuelle Registerwert ist  $a^{\frac{1}{2^x}}$ , er liegt immer näher bei 1. Entsprechend kann die Genauigkeit des Resultats bestimmt werden über den Wert von  $a^{\frac{1}{2^x}}$ , je nachdem welche Nachkommastelle noch beeinflusst wird. D so also beliebige Genauigkeit möglich ist, kennen wir  $x$  aus  $a^x = y$ .

Pseudocode zum entsprechenden Algorithmus:

```
Input: a,x,  a>0,0≤x<1
Output: ax
Variable: resultat, register, i, y
Setze resultat=1
      register=a
      i      =1
for i=1 bis i=MAXWERT
// auch bei abbrechender Dezimalbruchentwicklung
// bis MAXWERT Stellen zu berechnen, spart Zeit,
// weil sonst immer auf Gleichheit mit 0
// geprüft werden müsste
reg=sqrt(reg)
y=2x
Ist y≥1
      resultat=resultat*register
      y=y-1
setze x=y
i=i+1
endfor
```

In C-Code

```
#define MAXWERT 50
// Voraussetzung: a>0, 0≤x<1
// Output der Funktion pow ist ein Double ax
double pow(double a, double x)
{
    double resultat=1.0,
          register=a,
          y;
    int i;
    for (i=1;i≤MAXWERT;i++)
    {
        register=sqrt(register);
        y=2*x;
        if (y≥1)
        {
            resultat*=register;
            y-=1.0;
        }
        x=y;
    }
}
```

```

    }
    return (resultat);
}

```

Lösung von  $10^x = 2$ ? Mit diesem Verfahren wurden früher Logarithmentafeln per Hand errechnet. Dazu erstellt man zunächst eine Wertetabelle von  $\sqrt{10}$ ,  $\sqrt{\sqrt{10}}$  usw. Dann beginnt man mit Vergleichen:

$10^x = 2$ , umzuwandeln in  $10^{(0,b_1b_2\dots)_2}$

$10^0 = 1$ . Da  $1 < 2$ , muss  $x > 0$  sein.

$10^{0,1} = \sqrt{10} > 2$ . Dieser Wert ist zu groß. Also beginnt  $x$  mit  $0,0\dots$ . Was ist die zweite Stelle?

$10^{0,01}$  Man berechnet die Werte so, dass man immer unter 2 bleibt, aber immer näher an 2 herankommt.

**29.6. Effektive Berechnung der Quadratwurzel durch das Heron-Verfahren.** Dies ist nötig als Rechenoperation zu obigem Verfahren der Berechnung von  $10^x = 2$ . Algorithmus:

- (1) Gegeben sei  $a > 1$ ; ist  $a < 1$ , so setze man  $a := \frac{1}{a}$ , im wird also berechnet  $\sqrt{\frac{1}{a}}$ , und am Ende berechne man  $\sqrt{a} = \frac{1}{\sqrt{\frac{1}{a}}}$ .
- (2) Gesucht ist  $\sqrt{a}$ .
- (3) Man suche sich  $x_0^2 \approx a$  mit  $x_0^2 \leq a \leq (x_0 + 1)^2$ ,  $x_0 > 1$ .
- (4) Der abweichende Fehler ist  $|x_0^2 - a| \leq \varepsilon$ . Die gefundene Wurzel  $x_0 \approx \sqrt{a}$  ist gut, wenn  $\varepsilon \ll$ .
- (5) Setze  $x_1 = \frac{1}{2}(x_0 + \frac{a}{x_0})$ . (Warum?)  $x_1$  ist eine bessere Annäherung an die Quadratwurzel als  $x_0$ : Nun gilt:  $|x_1^2 - a| \leq \frac{\varepsilon^2}{4x_0^2} \leq \frac{1}{4}\varepsilon^2$ . Man wiederhole das Verfahren nun mit  $x_1$  statt  $x_0$ . Das Verfahren nähert sich sehr schnell an die Quadratwurzel an. Es heißt Heron-Verfahren. Es ist weit schneller als die Annäherung an die Quadratwurzel durch Verwendung des Mittelwertes als neuem Wert, wenn vorher ein  $x_0^2 > a$  und ein  $x_1^2 < a$  vorhanden waren. Das Heron-Verfahren ist ein wesentlicher Schritt, wenn der Computer (auch ein Taschenrechner!) eine Quadratwurzel berechnet; eine Iteration des Heron-Verfahrens ergibt jeweils eine Verdopplung der Nachkommastellen an Genauigkeit. Siehe Sourcecode zu bc, wo auch dieses Verfahren angewandt wird.

**29.7. Zahldarstellung im Computer: Datentypen int, long, float, double.** Ein Datentyp ist eine Information, wie eine Information dargestellt wird. int und long sind ganze Zahlen, float und double sind Fließkommazahlen. Im Folgenden Darstellung des Internationalen Standards zur Zahldarstellung.

**29.7.1. Darstellung ganzer Zahlen.** Gegeben sei eine ganze Zahl  $n$ . Wie kann ein Computer 16bit nutzen, um diese Zahl darzustellen? Das Vorzeichenbit  $b_{15}$  (Wert 1:  $-1$ ; Wert 0:  $+1$ ) wird als erstes Bit abgetrennt. Die folgenden 15bit  $b_{14} \dots b_0$  stellen eine 15stellige Binärzahl dar. Die größte mit 16 bit darstellbare Zahl ist also

$$0111111111111111_2 = 2^{14} + 2^{13} + \dots + 2^1 + 2^0 = 2^{15} - 1 = 32767$$

Jedoch ist die größte negative Zahl  $1111111111111111_2 = 2^{15} = -32768$ <sup>17</sup>. Dies liegt daran, dass sonst die Darstellung der 0 durch  $\pm 0$  redundant wäre; deshalb wird  $-1$  dargestellt als 1000000000000000, weshalb der Betrag der größten negativen Zahl 1 größer ist als der Betrag der größten darstellbaren positiven Zahl! Benennung der Datentypen:

- 1 Byte: Zahlbereich  $-128, \dots, +127$
- 2 Byte: Zahlbereich  $-32768, \dots, 32767$ . Namen: short, shortint
- 4 Byte: Zahlbereich ca.  $-2^{31}, \dots, +(2^{31} - 1)$ . namen: int, long
- 8 Byte: Zahlbereich ca.  $-2^{63}, \dots, +(2^{63} - 1)$ . Namen: long (in Java)

**29.7.2. Darstellung von Fließkommazahlen.** Darstellung des Kommas durch einen Punkt im Computerbereich, aufgrund des nationalen Standard in den USA. Festkommazahlen haben ein Komma an einer festen Stelle, z.B. Preise mit Pfennigen. Bei Fließkommazahlen steht jedoch das Komma an beliebiger Stelle. Die wissenschaftliche Darstellung von Fließkommazahlen im Dezimalsystem besteht aus einer ganzen Zahl zwischen 1 und 9, Nachkommastellen und einem Faktor  $10^x$ , zum Beispiel  $10.0 = 9,2654 \cdot 10^{-23}$ .

Ähnlich stellt der Computer Zahlen dar, indem er sie zerlegt in eine Mantisse und einen Exponenten zur Basis 2:

$$x = \pm m \cdot 2^e$$

$$10.0 = 1.25 \cdot 2^3 = (1.01)_2 \cdot 2^3$$

$e::$  Exponent  
 $m::$  Mantisse

<sup>17</sup>negative Zahlen werden im Computer eigentlich durch ihr Zweierkomplement, d.h. ein invertiertes Bitmuster der Binärzahl, dargestellt ( $1000000000000000_2 = -32768$  und  $1111111111111111_2 = -1$ ), aufgrund der Aufteilung aller in 16 Bit darstellbaren Binärzahlen in einen Zahlenkreis, in dem zu Beginn 0 und  $-1$  zusammentreffen und gegenüber, an der Stelle, wo das führende Bit wechselt,  $+32767$  und  $-32768$ . Hier wird zur Vereinfachung so getan, als würden positive und negative Zahlen im Rechner als normale Binärzahlen dargestellt.

Darstellung der Zahl 10.0 im Computer durch 64bit (Datentyp double) nach IEEE-Standard:

- das höchstwertige Bit ist das Vorzeichenbit
- die 11 folgenden Bit sind die Binärdarstellung des Exponenten. Um auch negative Exponenten darstellen zu können, steht hier die binäre Darstellung von  $E = e + 1023$ ,  $e$  ist der eigentliche Exponent. Man hätte theoretisch auch eine Art Darstellung als Integer mit Vorzeichenbit wählen können, jedoch ist die andere Darstellung zur Verarbeitung in der Hardware geeigneter.
- die 52 folgenden Bit sind die Binärdarstellung der Mantisse. Es sind die Nachkommastellen, dargestellt als Bitmuster im Binärsystem. Die führende Ziffer vor dem Komma ist im Binärsystem immer 1, kann also weggelassen werden.
- Also ist die Darstellung von 10.0 im Computer als float:

$$0|0000000011|010...0$$

Die letzten 50 Bits sind also Nullen. An dieser Darstellung wird klar, wie kompliziert die Addition von Fließkommazahlen ist, warum also die Ganzzahlenarithmetik am Computer so schnell ist. Man sollte sich beim Programmieren angewöhnen, nur Fließkommazahlen miteinander zu verrechnen:  $n = 7.4 + 2.0$  statt  $n = 7.4 + 2$ .

**Example 93.** Umrechnung einer Fließkommazahl in Binärdarstellung nach IEEE-Standard double

**Example.**

- Gegeben ist  $x = -419,8125$ .
- Zuerst wandelt man in Binärformat um:  $x_2 = -110100011,1101_2$ .
- Dann verschiebt man das Komma um  $n$  Stellen hinter die führende 1 und gleicht dies durch Multiplikation mit  $2^n$  aus. Es ergibt sich eine Darstellung wie  $x_2 = -1,101000111101_2 \cdot 2^8$ .
- Nun kann die Darstellung als double abgelesen werden:
  - Vorzeichenbit: Wert 1 für negative Zahlen, 0 für positive Zahlen.
  - Exponent: die Binärdarstellung des Exponenten, addiert mit 1023.
  - Mantisse: die Nachkommastellen, aufgefüllt mit Nullen bis auf 52 Bit.

## 29.8. Aufbau der Zahlen.

$$\mathbb{N} \subset \mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R} \subset \mathbb{C}$$

29.8.1. *Rationale Zahlen.* Rationale Zahlen sind das Verhältnis von zwei ganzen Zahlen zueinander; man kann dieses Verhältnis aufschreiben als  $\frac{p}{q}$  oder auch beliebig anders, z.B. als  $(1, 2)$ . Durch Äquivalenzumformungen kann geschrieben werden  $(1, 2) = (2, 4)$ , allgemein bekannt als Kürzen. Zwei Zahlen können soweit gegeneinander gekürzt werden, bis sie teilerfremd sind. In dieser Form wird ein Bruch aufgeschrieben.

Addition zweier rationaler Zahlen als:

$$\frac{p}{q} + \frac{r}{s} = \frac{p \cdot s}{q \cdot s} + \frac{r \cdot q}{s \cdot q} = \frac{ps + rq}{qs}$$

Darstellung rationaler Zahlen als Dezimalbruch. Die Darstellung rationaler Zahlen als  $\frac{p}{q}$  wurde schon von den Griechen angewandt, die Darstellung als Dezimalbruch kam jedoch erst im 14. Jahrhundert auf:  $\frac{1}{4} = 0,25$ . Durch die Erweiterung auf nicht-abbrechende Dezimalbrüche kam der Begriff der reellen Zahlen auf; verboten ist allein eine Zahlentwicklung  $n_1, n_2 \dots n_3 \overline{9}$ . Das wiederum führte nach 300 Jahren Nachdenken zum Begriff des Grenzwertes als Ziel (formuliert von Weierstraß um 1900), an das sich ein nicht-abbrechender Dezimalbruch annähert.

Wie erkenne ich die ganzen Zahlen in den rationalen Zahlen? Die rationalen Zahlen enthält eine Teilmenge (sei  $\overline{\mathbb{Z}} \subset \mathbb{R}$ ) mit Elementen der Gestalt  $\frac{n}{1}$ ,  $n \in \mathbb{N}$ . Diese Elemente werden mit den Elementen aus  $\mathbb{Z}$  identifiziert, denn es sind dieselben Zahlen in anderer Darstellung.

Wie kann man nun die rationalen Zahlen in der Menge der reellen Zahlen  $\mathbb{R}$ ? Die Umwandlung rationaler Zahlen in Dezimalbrüche geschieht durch sukzessive Division mit Rest, wobei  $10 \cdot \text{Rest}$  als neuer Dividend verwendet wird. Bei einer Division  $p : q$  gibt es  $q$  Möglichkeiten für den Rest, nämlich  $[0, \dots, q-1]$ . Taucht bei der Division 0 als Rest auf, so endet der Dezimalbruch. Taucht ein Rest zum zweitenmal auf, so beginnt die Dezimalbruchentwicklung periodisch von neuem. Die Division mit Rest ist:

$$\frac{p}{q} = n + \frac{r}{q} \quad 0 \leq r < q$$

Also führt eine rationale Zahl zu einer periodischen oder einer abbrechenden Dezimalbruchentwicklung.

Umwandlung einer abbrechenden Dezimalbruchentwicklung in eine rationale Zahl:

$$0,125 = \frac{0,125}{1} = \frac{0,125 \cdot 1000}{1 \cdot 1000} = \frac{125}{1000}$$

Umwandlung einer periodischen Dezimalbruchentwicklung (Periodenlänge  $a$ ) in eine rationale Zahl:

$$\begin{aligned} 10^a \cdot x &= \text{Periode, Periode Periode} \dots \\ x &= 0, \text{Periode Periode} \dots \\ (10^a - 1) \cdot x &= \text{Periode} \\ x &= \frac{\text{Periode}}{10^a - 1} \end{aligned}$$

Oder bei periodischen Dezimalbrüchen mit Vorperiode an einem Beispiel:

$$\begin{aligned} x &= 0,1\overline{3} \\ 100x &= 13,\overline{3} \\ 10x &= 1,\overline{3} \\ (100 - 10)x = 90x &= 12 \\ x &= \frac{12}{90} = \frac{2}{15} \end{aligned}$$

29.8.2. *Reelle Zahlen.* Bezeichnung der Zahlenmenge mit  $\mathbb{R}$ , der Elemente mit  $a, b, c, \dots$  oder auch mit beliebigen anderen Symbolen. Eine Menge  $R$ , für die alle folgenden Gesetzgeboten, heißt ein »Körper«.  $\mathbb{R}, \mathbb{Q}$  sind Körper, wobei  $\mathbb{Q} \subset \mathbb{R}$  gilt: die rationalen Zahlen sind ein Unterkörper der reellen Zahlen. Es kann bewiesen werden, dass  $\mathbb{Q}, \mathbb{R}, \mathbb{C}$  unter bestimmten Voraussetzungen die einzigen Körper im Zahlensystem sind.

- Es sind die Verknüpfungen  $a + b, a - b, a \cdot b = a * b = ab, \frac{a}{b}$  für  $b \neq 0$  definiert, deren Ergebnisse stets auch wieder reelle Zahlen sind.
- Kommutativgesetz:
  - der Addition:  $a + b = b + a$
  - der Multiplikation:  $a \cdot b = b \cdot a$
- Assoziativgesetz: Gilt das Assoziativgesetz für eine beliebige Verknüpfung  $\circ$ , so kann ein Ausdruck wie  $a \circ b \circ c \circ d$  beliebig geklammert werden, ohne das Ergebnis zu verändern. Gilt zusätzlich das Kommutativgesetz, so ist das Ergebnis auch unabhängig von der Klammerung.
  - der Addition:  $(a + b) + c = a + (b + c)$
  - der Multiplikation:  $(a \cdot b) \cdot c = a \cdot (b \cdot c)$
- Distributivgesetz:  $(a + b) \cdot c = ac + bc$
- Neutrales Element
  - der Addition: es gibt  $0 \in \mathbb{R}$ , für das gilt:  $a + 0 = a$
  - der Multiplikation: es gibt  $1 \in \mathbb{R}$ , für das gilt:  $a \cdot 1 = a$
- Inverses Element
  - der Addition: zu jedem  $a \in \mathbb{R}$  gibt es ein  $-a \in \mathbb{R}$ , so dass  $a + (-a) = 0$
  - der Multiplikation: zu jedem  $a \in \mathbb{R}$  gibt es ein  $a^{-1} = \frac{1}{a} \in \mathbb{R}$ , so dass  $a \cdot \frac{1}{a} = 1$ .

Anordnung der reellen Zahlen. Anordnung der reellen Zahlen: auf einer Zahlengerade liegen alle reellen Zahlen.  $a < b$  bedeutet:  $a$  liegt auf dem Zahlenstrahl links von  $b$ ;  $a > b$  bedeutet:  $a$  liegt rechts von  $b$ .

Teilmengen von  $\mathbb{R}$ :

$$\mathbb{R}_+ = \{x \in \mathbb{R} | x \geq 0\}$$

,

$$\mathbb{R}_- = \{x \in \mathbb{R} | x \leq 0\}$$

$$\mathbb{R}_+ \cap \mathbb{R}_- = \{0\}$$

$$\mathbb{R}^* = \mathbb{R} \setminus \{0\}$$

Die positiven Zahlen:  $\mathbb{R}_+^* = \mathbb{R}^* \cap \mathbb{R}_+$

Die negativen Zahlen:  $\mathbb{R}_-^* = \mathbb{R}^* \cap \mathbb{R}_-$

Warum ist  $-1 \cdot -1 = +1$ ?

$$\begin{aligned} 1 + (-1) &= 0 \\ (1 + (-1)) \cdot (-1) &= 0 \cdot (-1) = 0 \\ 1 \cdot (-1) + (-1) \cdot (-1) &= 0 \\ -1 + (-1) \cdot (-1) &= 0 \\ (-1) \cdot (-1) &= 1 \end{aligned}$$

Der Betrag einer Zahl. Definition des Betrages einer Zahl  $a \in \mathbb{R}$ :

•

$$|a| = \begin{cases} a & a \geq 0 \\ -a & a < 0 \end{cases}$$

- besser:  $|a| = \sqrt{a^2} \Leftrightarrow |a|^2 = a^2$ . Mit  $\sqrt{a}$  meint man stets die positive Wurzel aus einer Zahl!

Was ist nun  $(|a| + |b|)^2$ ?

$$\begin{aligned} (|a| + |b|)^2 &= |a|^2 + 2 \cdot |a| \cdot |b| + |b|^2 \\ &= |a|^2 + 2 \cdot |ab| + |b|^2 \end{aligned}$$

Es ist der Beweis nötig, dass  $|a| \cdot |b| = |ab|$ .

Es gilt stets:  $|a| \geq a$ ;  $|a| + |b| \geq |a + b|$ . Beweis?

Intervalle. Das abgeschlossene Intervall ist:

$$[a, b] = \{x \in \mathbb{R} \mid a \leq x \leq b\}$$

Das offene Intervall ist:

$$]a, b[ = \{x \in \mathbb{R} \mid a < x < b\}$$

Die halboffenen Intervalle sind:

$$]a, b] = \{x \in \mathbb{R} \mid a < x \leq b\}$$

$$[a, b[ = \{x \in \mathbb{R} \mid a \leq x < b\}$$

Gilt  $|a - b| = |b - a|$ ?

$$\begin{aligned} (-1)(a - b) &= (b - a) \\ \Rightarrow |(-1)(a - b)| &= |b - a| \\ \Leftrightarrow |a - b| &= |b - a| \end{aligned}$$

qed.

Der Betrag  $|a - b|$  kann interpretiert werden als Abstand der Punkte  $a, b$  auf der Zahlengeraden bzw. die Länge der Strecke zwischen  $a$  und  $b$  auf der Zahlengeraden.

Angabe der Zahlen, die bis zu einem maximalen Abstand von 3 um eine Zahl  $a$  liegen:  $[a - 3, a + 3]$  oder über den Abstand:  $|a - x| \leq 3$ . Also gilt:

$$[a - 3, a + 3] = \{x \in \mathbb{R} \mid |a - x| \leq 3\}$$

$$]a - 3, a + 3[ = \{x \in \mathbb{R} \mid |a - x| < 3\}$$

Damit können Umgebungen von  $a$  angegeben werden:

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

**Example 94.** Gilt das Assoziativgesetz für die Subtraktion?

**Example.** Berechnet man  $7 - 4 - 5$  von links, so ist das Ergebnis  $-2$ . Berechnet man den Ausdruck von rechts, so ist das Ergebnis  $8$ . Die definierten Verknüpfungen der reellen Zahlen verknüpfen stets nur zwei reelle Zahlen. Deshalb ist ein Ausdruck  $7 - 4 - 5$  mathematisch sinnlos; er kann aber unterschiedlich geklammert werden, wodurch jeweils wieder Verknüpfungen von zwei Elementen entstehen:

$$(7 - 4) - 5 = -2 \neq 8 = 7 - (4 - 5)$$

Die Subtraktion erfüllt also nicht das Assoziativgesetz.

**Example 95.** Gilt das Assoziativgesetz für das Potenzieren?

**Example.** Nein, denn

$$(2^3)^4 = 2^{12} \neq 2^{21} = 2^{(3^4)}$$

Also macht ein Ausdruck  $2^{3^4}$  ohne Klammerung mathematisch keinen Sinn.

Rechnen mit Ungleichungen in  $\mathbb{R}$ . Bei Addition einer beliebigen reellen Zahl  $c$  bleibt die Ungleichung erhalten:

$$\begin{aligned} a &< b \\ a + c &< b + c \end{aligned}$$

Bei Multiplikation mit  $c \in \mathbb{R}_+^*$  bleibt die Ungleichung erhalten:

$$\begin{aligned} a &< b \\ ac &< bc \end{aligned}$$

Bei Multiplikation mit  $c \in \mathbb{R}_-^*$  dreht sich die Ungleichung um; es muss also bei Multiplikation mit einer Variablen ggf. eine Fallunterscheidung durchgeführt werden:

$$\begin{aligned} a &< b \\ ac &> bc \end{aligned}$$

Bei Multiplikation mit  $c = 0$  entsteht  $0 = 0$ .

29.8.3. *Natürliche Zahlen*. Was die natürlichen Zahlen sind, ist letztlich unbekannt. Sie werden auch nicht in Frage gestellt. Es können jedoch Eigenschaften angegeben werden, die erfüllt sein müssen, damit eine Menge die Menge  $\mathbb{N}$  ist, die also die natürlichen Zahlen charakterisieren:

Prinzip der vollständigen Induktion: Gegeben sei  $M \subset \mathbb{N}$  mit

- (1)  $1 \in M$
- (2) liegt  $k \in M$ , so zeige man dass  $k + 1 \in M$ . Für diesen Beweis gibt es kein Standardverfahren. Diese Formulierung  $n \in M \Rightarrow n + 1 \in M$  wird manchmal auch äquivalent formuliert als:  $1, 2, 3, \dots, n \in M \Rightarrow n + 1 \in M$

so gilt  $M = \mathbb{N}$ .

Vollständige Induktion zeigt vollständig die Identität von Zahlenmengen. Unvollständige Induktion dagegen wäre: da es auch große Primzahlen gibt, sind alle Zahlen Primzahlen.

**Example 96.** Beweis, dass  $M$  die Menge  $\mathbb{N}$  ist

**Example.** Behauptung:

$$\begin{aligned} M &= \mathbb{N} \\ F(n) &= \mathbb{N} \\ \frac{n(n+1)}{2} &= 1 + 2 + 3 + \dots + n \end{aligned}$$

Liegt  $1 \in M$ ? Ja, denn  $1 = \frac{1(1+1)}{2}$ .

Liegt  $k \in M$ ?

$$1 + 2 + \dots + k = \frac{k(k+1)}{2}$$

Liegt  $k + 1 \in M$ ?

$$(24) \quad 1 + 2 + \dots + k + (k + 1) = \frac{(k + 1)(k + 2)}{2}$$

Wenn nun ausgehend von der postulierten Behauptung  $1 + 2 + \dots + k = \frac{k(k+1)}{2}$  gezeigt werden kann, dass auch 24 gilt, so wurde gezeigt, dass  $M = \mathbb{N}$ . Also:

$$\begin{aligned} 1 + 2 + \dots + k &= \frac{k(k+1)}{2} \quad | + (k + 1) \\ 1 + 2 + \dots + k + (k + 1) &= \frac{k(k+1) + 2(k+1)}{2} \end{aligned}$$

Damit gilt die postulierte Formeln für alle natürlichen Zahlen. Bewiesen wurde dies über das Prinzip der vollständigen Induktion, über die eigentlich alle Beweise zu natürlichen Zahlen gehen. Dieses Prinzip wurde von dem Italiener Peano formuliert.

**Example 97.** Vollständige Induktion

**Example.** Vermutung:

$$Q : 1^2 + 2^2 + \dots + n^2 = \frac{n(n+1)(2n+1)}{6}$$

Es gilt  $1 \in Q$ :

$$Q(1) : 1^2 = \frac{1 \cdot 2 \cdot 3}{6}$$



Unter Voraussetzung der Behauptung  $Q$  soll gelten  $k+1 \in Q$ . Um zu dieser Behauptung zu gelangen, addiert man auf beiden Seiten  $(k+1)^2$ , um so die Folgerung  $k \in Q \Rightarrow k+1 \in Q$  zu beweisen:

$$\begin{aligned} 1^2 + 2^2 + \dots + k^2 + (k+1)^2 &= \frac{k(k+1)(2k+1)}{6} + (k+1)^2 \\ &= (k+1) \cdot \left[ \frac{k \cdot (2k+1)}{6} + (k+1) \right] \\ &= (k+1) \cdot \frac{2k^2 + 7k + 6}{6} \\ &= (k+1) \cdot \frac{(k+2)(2k+3)}{6} \\ &= \frac{(k+1) \cdot ((k+1)+1) \cdot (2(k+1)+1)}{6} \end{aligned}$$

Da die gewünschte Zielformel  $((k+1)$  eingesetzt für  $n$  in  $\frac{n(n+1)(2n+1)}{6}$ ) entstand, wurde nach der vollständigen Induktion bewiesen, dass obige Vermutung für jedes  $n \in \mathbb{N}$  gilt.

### Example 98. Vollständige Induktion

**Example.** Die vollständige Induktion beginnt stets mit einer Formel (z.B.  $n^2 = -35$ ) und der Behauptung, die Formel sei richtig für alle  $n \in \mathbb{N}$ . Behauptung hier: für  $a \geq -1$ ,  $a \in \mathbb{R}$  gilt  $(1+a)^n \geq 1+n \cdot a$ . Beweis: Sei  $M \subset \mathbb{N}$  die Teilmenge, für die diese Formel richtig ist.

- (1) zeige dass  $1 \in M$ :

$$\begin{aligned} (1+a)^1 &\geq 1+1 \cdot a \\ \Leftrightarrow 1+a &\geq 1+a \quad (w) \end{aligned}$$

- (2) Sei  $n \in M$ . Das heißt:  $(1+a)^n \geq 1+n \cdot a$  für  $a \geq -1$ . Nun zeige, logisch, dass aus dieser Formel die Zielformel bei der vollständigen Induktion folgt, nämlich:

$$(1+a)^{n+1} \geq 1+(n+1) \cdot a \text{ für } a \geq -1$$

Dazu multipliziere man mit  $1+a$ ; für  $1+a \geq 0 \Leftrightarrow a \geq -1$  gilt dann, also in Übereinstimmung mit der Ausgangsbedingung:

$$\begin{aligned} (1+a)^n \cdot (1+a) &\geq (1+n \cdot a) \cdot (1+a) \\ \Leftrightarrow (1+a)^{n+1} &\geq 1+(n+1) \cdot a + n \cdot a^2 \end{aligned}$$

Da  $n \cdot a^2 \geq 0$ , gilt auch:

$$(1+a)^{n+1} \geq 1+(n+1) \cdot a$$

Dies ist die Zielformel der vollständigen Induktion. Also wurde bewiesen, dass  $M = \mathbb{N}$ .

### Example 99. Vollständige Induktion: Jedes $n \in \mathbb{N}$ ist Produkt von Primzahlen.

#### Example.

- (1) Es gilt  $1 \in M$
- (2) Weiter gilt der Satz für  $n=2$ ,  $n=3$ , also sei er für  $n$  richtig. (andere Formulierung der vollständigen Induktion).
- (3) Gilt der Satz nun auch für  $n+1$ ? Ja, denn entweder ist  $n+1$  eine Primzahl, oder es lässt sich zerlegen in  $n+1 = p \cdot q$  mit  $p, q < n$ . Die Zerlegung ist rekursiv weiter möglich, bis eine vollständige Zerlegung von  $n+1$  in Primfaktoren besteht. Also gilt der Satz für  $n+1$  und damit gemäß der vollständigen Induktion für jedes  $n \in \mathbb{N}$ .

## 30. PROF. LÖFFLER: POLYNOME

Der Ausdruck  $ax^n$  mit  $a \in \mathbb{R}$ ,  $n \in \mathbb{N} > 0$  ist eine Rechenvorschrift die besagt, was man mit einer Zahl  $x$  zu tun hat. Man nennt diesen Ausdruck ein Monom mit dem Koeffizienten  $a$ , der Unbestimmten  $x$  und dem Grad  $n$ . Polynome nun sind Summen von Monomen (unter Verwendung von  $a_0x^0 = a_0 \cdot 1$ ):

$$p(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x^1 + a_0$$

Dabei werden die Monome wie hier nach absteigendem oder auch nach aufsteigendem Grad geordnet. Polynome  $p(x)$  können als Rechenvorschrift für  $x$  interpretiert werden (und sind dann die wichtigste Sorte Rechenvorschriften in der Mathematik) oder als formaler Ausdruck. Der Grad eines Polynoms ist gleich dem höchsten enthaltenen Grad eines Monoms, das einen Koeffizienten  $a \neq 0$  hat, geschrieben mit einer »mathematischen

Bewertung« als  $|p(x)| = n$ . Die Koeffizienten sind  $a_i \in \mathbb{R}$ , es können jedoch Einschränkungen gemacht werden. Das Bildungsgesetz der Polynome ist unter Verwendung des Summenzeichens:

$$p(x) = \sum_{i=0}^n a_i \cdot x^i$$

Addition zweier Polynome:

$$r(x) = p(x) + q(x) = \sum_{i=0}^n a_i \cdot x^i + \sum_{i=0}^m b_i \cdot x^i$$

für  $n = m$  ist die Notation einfach:

$$r(x) = p(x) + q(x) = \sum_{i=0}^n (a_i + b_i) \cdot x^i$$

Trick: Wenn  $n \neq m$ , so wird der Grad angeglichen durch Hinzufügen von Koeffizienten, die 0 sind, zum Polynom mit dem niedrigeren Grad.

Multiplikation zweier Polynome: Unter Verwendung des Distributivgesetzes und anschließendem Zusammenfassen.

Neben Monomen in einer Unbestimmten  $ax^n$  gibt es auch Monome in zwei Unbestimmten  $k \cdot x^a \cdot y^b$ . Was ist der Grad solcher Monome? Er ist  $a + b$  für den Fall  $x = y$ . Ein Polynom mit zwei Unbestimmten wird geschrieben als:

$$\begin{aligned} b &= m \\ a &= n \\ \sum_{a=0}^n \sum_{b=0}^m k_{ab} \cdot x^a \cdot y^b \end{aligned}$$

Der Grad eines Monoms in zwei Unbestimmten ist immer die Summe der Grade der beiden Unbestimmten. Somit sind alle möglichen Monome in zwei Unbestimmten vom Grad 3 mit dem Koeffizienten  $k = 1$ :

$$x^3, x^2y, xy^2, y^3$$

**30.1. Auswerten eines Polynoms an der Stelle  $\beta$ .** Dies ist ein Standardproblem in der Mathematik. Gegeben ist das Polynom  $p(x)$  vom Grad  $n \in \mathbb{R}$ , gesucht  $p(\beta)$  für  $\beta \in \mathbb{R}$ .

**Example 100.** »naives« Auswerten eines Polynoms

**Example.** Gegeben ist  $p(x) = 3x^3 - 7x^2 + 12x - 5$ . Was ist  $p(2)$ ? Für die Berechnung  $p(2) = 3 \cdot 2^3 - 7 \cdot 2^2 + 12 \cdot 2 - 5 = 15$  benötigt man insgesamt 5 Multiplikationen und 3 Additionen.

Ausgehend von Beispiel 100 benötigt man für dieses naive Auswerten eines Polynoms vom Grad  $n$   $n - 1$  Multiplikationen zum Berechnen von  $x^n$  bis  $x^2$ ,  $n$  Multiplikationen zum Multiplizieren mit den Koeffizienten und zusätzlich  $n$  Additionen. Der Rechenaufwand lässt sich durch Verwenden des Horner-Schemas reduzieren. Dazu schreibt man:

$$\begin{aligned} p(x) &= 3x^3 - 7x^2 + 12x - 5 \\ &= ((3x - 7) \cdot x + 12) \cdot x - 5 \end{aligned}$$

Das Polynom lässt sich also allgemein als Ausdrücke der Form  $A = A \cdot x + b$  zergliedern. In diesem Schema benötigt man zum Auswerten eines Polynoms vom Grad  $n$  nur  $n$  Multiplikationen und  $n$  Additionen.

**Example 101.** Auswerten eines Polynoms nach dem Horner-Schema

**Example.** Gegeben ist

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

Im Schema für  $p(2)$ :

|   |   |    |   |    |    |    |    |     |
|---|---|----|---|----|----|----|----|-----|
|   | 3 | -6 | 5 | -4 | 3  | -2 | 0  | 1   |
| 2 |   | 6  | 0 | 10 | 12 | 30 | 56 | 112 |
|   | 3 | 0  | 5 | 6  | 15 | 28 | 56 | 113 |

In der ersten Zeile stehen die Koeffizienten; in der zweiten Zeile stehen die Produkte aus dem Ergebnis der letzten Spalte und dem Wert 2 für die auszuwertende Stelle; in der dritten Zeile steht immer das Ergebnis der Summe aus den beiden Spalten darüber. Das Endergebnis steht also unten rechts. Umsetzung des Horner-Schemas in eine C-Routine:

```
// ein Array von Koeffizienten, vollständig und absteigend
// geordnet:
double coeff[]={7.0,0.0,0.0,2.0,1.0,-3.0,8.0};
// entspricht coeff[7]= usw., jedoch unnötig (Elementzahl)
// Definition des Horner-Schemas
// aus grad berechnet die Routine die Anzahl der
// Koeffizienten zu grad+1
// wert: die Stelle, an der das Polynom ausgewertet werden soll
double horner(double wert,double coeff[],int grad)
// nicht definierter Rückgabewert ist immer int
{
    // res: das Resultat, das Ergebnis der Gleichung
    double res=0.0; // Resultat
    int i;
    for (i=0; i<=grad; i++)
    {
        // Ansprechen des Arrays wie in PASCAL, vgl. aber unten
        res=res*wert+coeff[i]
    }
    return(res);
}
```

Nun gibt es in C zur effizienteren Behandlung des Zugriffs auf Arrays und andere Zwecke zu jedem Datentyp den Datentyp Zeiger, z.B. `double *coeff`, der den Wert der Speicheradresse einer Variablen enthält. Das nächste Elements im Array kann dann unter Verwendung des Zeigers `coeff` mit `coeff[coeff+1]` angesprochen werden; durch +1 wird dabei die nächste Adresse angesprochen. Diese Art der Auswertung ist weit schneller, als stets eine vollständig neue Adresse für ein Element im Array berechnen zu müssen.

In C wird bei Übergabe eines Arrays an eine Routine nur dessen Basisadresse `&coeff[0]` übergeben. Beim Zugriff auf das Array wird dann die Adresse eines Elements  $n$  berechnet über `&coeff[0] + sizeof(double) * n`. Nun gibt es in C zu jedem Datentyp einen Pointer-Datentyp, der eine Adresse enthält, an der ein Wert mit dem entsprechenden Datentyp steht. Angegeben als: `double *coeff`, so dass `coeff` vom Datentyp `double *` ist. Beim Aufruf wäre `coeff` die Adresse selbst und `*coeff` der Inhalt der Speicherzelle, auf die die Adresse zeigt. So kann eine effizientere Variante des Programms geschrieben werden:

```
double horner(double wert,double coeff[],int grad)
// nicht definierter Rückgabewert ist immer int
{
    // res: das Resultat, das Ergebnis der Gleichung
    double res=0.0, *dblptr=&coeff[0];
    int i=grad+1;
    // entspricht for (; i>0;i--), d.h. Test auf FALSE:
    for (; i; i--)
    {
        // Inhalt *dblptr auslesen, danach dblptr erhöhen;
        // es wird um 8 Byte erhöht, da eine Einheit von
        // double *dblptr 8 Byte ist.
        res=res*wert+ *dblptr++;
        // dies ist die gute Art, in C Arrays anzusprechen
    }
    return(res);
}
```

Noch kürzere Variante unter Verwendung einer äquivalenten WHILE-Schleife:

```
double horner(double coeff[], int n, double wert)
{
    double res=0.0, *dblptr=&coeff[0];
    int laufindex=n+1;
    while (laufindex--)
    {
        res=res*wert+ *dblptr++;
    }
    return(res);
}
```

**30.2. Nullstellen von Polynomen.** Dies ist ein Standardproblem in der Mathematik. Gegeben ist das Polynom  $p(x)$  vom Grad  $n \in \mathbb{R}$ , gesucht ist ein  $\alpha \in \mathbb{R}$ , für das  $p(\alpha) = 0$  gilt.

**30.2.1. Nullstellen von Polynomen  $p(x)$  vom Grad  $n = 0$ .** Polynome vom Grad  $n = 0$  haben keine Nullstellen, außer  $p(x) = 0$ , das Nullstellen für jede Zahl hat.

**30.2.2. Nullstellen bei Polynomen  $p(x)$  vom Grad  $n = 1$ .**

$$\begin{aligned} ax + b &= 0 \\ ax &= -b \\ x &= -\frac{b}{a} \end{aligned}$$

Da  $n = 1$ , ist  $a \neq 0$  es kann also ohne Einschränkung durch  $a$  dividiert werden. Ein Polynom vom Grad  $n = 1$  hat also in jedem Fall eine Nullstelle.

**30.2.3. Nullstellen bei Polynomen  $p(x)$  vom Grad  $n = 2$ .** Die binomischen Formeln sind herzuleiten aus dem Distributivgesetz:

$$\begin{aligned} (a + b)^2 &= a^2 + 2ab + b^2 \\ (a - b)^2 &= a^2 - 2ab + b^2 \\ (a + b)(a - b) &= a^2 - b^2 \end{aligned}$$

Jedes Polynom vom Grad  $n = 2$  kann nun durch Addition und anschließende Subtraktion einer Zahl zu einem vollständigen Quadrat ergänzt werden, das nach einer binomischen Formel in Linearfaktoren umgeschrieben werden kann:

$$\begin{aligned} x^2 + px + q &= 0 \\ x^2 + px &= -q \\ \Leftrightarrow x^2 + px + \left(\frac{p}{2}\right)^2 - \left(\frac{p}{2}\right)^2 &= -q \\ \Leftrightarrow \left(x + \frac{p}{2}\right)^2 &= -q + \frac{p^2}{4} \end{aligned}$$

Aus dieser Form, die durch die sog. »quadratische Ergänzung« entstand, können nun die Nullstellen durch Wurzelziehen bestimmt werden, gleichzeitig ergibt sich die  $pq$ -Formel als Lösungsformel zur Bestimmung der Nullstellen von quadratischen Gleichungen, die in normierter Form  $x^2 + px + q$  gegeben sind:

$$\begin{aligned} x + \frac{p}{2} &= \pm \sqrt{\frac{p^2}{4} - q} \\ \Leftrightarrow x_{1,2} &= -\frac{p}{2} \pm \sqrt{\frac{p^2}{4} - q} \end{aligned}$$

Für quadratische Gleichungen in Normalform  $ax^2 + bx + c$  kann stattdessen auch die  $abc$ -Lösungsformel angewandt werden:

$$x_{1,2} = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a}$$

Eine quadratische Gleichung hat im allgemeinen zwei Lösungen  $x_{1,2}$ , sie kann jedoch auch nur eine Lösung haben (wie  $x^2 + 2x + 1 = 0 \Leftrightarrow x_1 = -1 \pm \sqrt{1 - 1}$ ) oder keine, wenn die Diskriminante, d.i. das Argument der Wurzel,  $D < 0$  ist.

**30.2.4. Nullstellen bei Polynomen  $p(x)$  vom Grad  $n \geq 3$ .** Polynome vom Grad  $n = 3$  und  $n = 4$  konnten bereits vor einigen hundert Jahren mit entsprechenden Lösungsformeln gelöst werden. Dazu gibt es z.B. die Cardano'sche Formel. Nach langen vergeblichen Versuchen, die Nullstellen von Polynomen vom Grad  $n \geq 5$  allgemein zu finden, konnte durch Galois und Abel gezeigt werden, dass hierfür allgemein keine Lösungsformel angegeben werden kann.

### 30.3. Die allgemeine binomische Formel.

30.3.1. *Darstellung.* Es ist berechenbar durch einfaches Ausmultiplizieren:

$$\begin{aligned}(a+b)^0 &= 1 \\ (a+b)^1 &= a+b \\ (a+b)^2 &= a^2+2ab+b^2 \\ &\dots\end{aligned}$$

Was ist nun das allgemeine Schema in dieser binomischen Formel? Es kann vermutet werden, dass das Ergebnis einer binomischen Formel  $n$ -ten Grades eine Summe aller Monome in zwei Unbestimmten  $a \cdot b^y$  mit dem Grad  $x+y=n$  ist. Was sind nun die Koeffizienten in dieser Summe von Monomen? Festlegung der Bezeichnungen für die Koeffizienten in allgemeiner Form:

$$\begin{aligned}(a+b)^n &= \sum_{i=0}^n \beta_i^n \cdot a^i b^{n-i} \\ &= \beta_0^n a^n + \beta_1^n a^{n-1} b + \beta_2^n a^{n-2} b^2 + \dots + \beta_{n-1}^n a b^{n-1} + \beta_n^n b^n\end{aligned}$$

In der Summenschreibweise ist dabei der Index  $n$  zu  $\beta$  der Grad des Polynoms, also  $n$  zu  $(a+b)^n$ , gleichzeitig die Ebene im Pascalschen Dreieck, und der Index  $i$  zu  $\beta$  die Stelle des Monoms in zwei Unbestimmten in der Summe, bei absteigender Ordnung nach den Exponenten von  $a^i$ .

Vermutet wird außer obiger Formel, dass die Koeffizienten nach dem Pascalschen Dreieck wie folgt berechnet werden können (allgemeine Regel zur Koeffizientenberechnung nach dem Pascalschen Dreieck):

$$(25) \quad \beta_{k+1}^{n+1} = \beta_k^n + \beta_{k+1}^n$$

mit  $\beta_0^n = 1$  und  $\beta_n^n = 1$ . Beispiele im Pascalschen Dreieck:  $\beta_3^5 = \beta_2^4 + \beta_3^4 = 3 + 3 = 6$ . Die Koeffizientenberechnung erfolgt also rekursiv.

30.3.2. *Beweis nach der vollständigen Induktion.* Definitionen

- Definition der Fakultät:  $n! = 1 \cdot 2 \cdot 3 \cdot \dots \cdot n$  und  $0! = 1$ .
- Außerdem benötigt man das Symbol » $n$  über  $k$ «:

$$\binom{n}{k} = \frac{n!}{k! \cdot (n-k)!}$$

für  $0 \leq k \leq n$ . Diese Formel kann stets gekürzt werden zu:

$$\binom{n}{k} = \frac{n \cdot (n-1) \cdot \dots \cdot (n-k+1)}{1 \cdot 2 \cdot 3 \cdot \dots \cdot k}$$

Weiter soll per Definition gelten:

$$\binom{n}{0} = \binom{n}{n} = 1$$

Bewiesen werden muss nach der vollständigen Induktion, dass die allgemeine binomische Formel und die Formel zur Berechnung der Koeffizienten nach dem Pascalschen Dreieck allgemein für jeden Grad  $n \in \mathbb{N}$  gilt.

Stichprobenweise Interpretation der Koeffizienten  $\beta_k^n$  als  $\binom{n}{k}$  zeigt, dass die richtigen Koeffizienten als Ergebnisse entstehen. Aus der Struktur des Pascalschen Dreiecks (die Koeffizientenreihen beginnen von vorne und hinten gleich) vermutet man nun:

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

Dies konnte also auch allgemein gezeigt werden. Die allgemeine Regel zur Koeffizientenberechnung nach dem Pascalschen Dreieck (Gleichung 25) wird nun, in der Schreibweise » $n$  über  $k$ « vermutet als:

$$\begin{aligned}\binom{n+1}{k} &= \binom{n}{k-1} + \binom{n}{k} \\ \binom{n}{k} &= \binom{n-1}{k-1} + \binom{n-1}{k}\end{aligned}$$

Der Aufbau einer Struktur wie die des Pascalschen Dreiecks nach dieser Formel liefert dieselben Koeffizienten; deshalb wird also vermutet für die Koeffizienten in der allgemeinen binomischen Formel:

$$\beta_k^n = \binom{n}{k}$$

Ist diese Behauptung richtig, so kann die allgemeine binomische Formel geschrieben werden als:

$$\begin{aligned}(a+b)^n &= \binom{n}{n} a^n + \binom{n}{n-1} a^{n-1} b + \binom{n}{n-2} a^{n-2} b^2 + \dots + \binom{n}{1} a b^{n-1} + \binom{n}{0} b^n \\ &= \binom{n}{0} a^n + \binom{n}{1} a^{n-1} b + \binom{n}{2} a^{n-2} b^2 + \dots + \binom{n}{n-1} a b^{n-1} + \binom{n}{n} b^n\end{aligned}$$

Dabei wurden die Beziehungen

$$\binom{n}{0} = \binom{n}{n} = 1$$

verwendet. In Summenschreibweise:

$$(a+b)^n = \sum_{i=0}^n \binom{n}{i} a^i b^{n-i}$$

Diese Formel soll nun bewiesen werden.

Sie gilt auf jeden Fall für  $n=0$ , denn  $x^0 = 1$ .

Sie gilt auch für den Fall  $n=1$ :

$$\begin{aligned}(a+b)^1 &= \sum_{i=0}^1 \binom{1}{i} a^i b^{1-i} \\ &= \binom{1}{0} a^0 b^1 + \binom{1}{1} a^1 b^0 \\ &= 1 \cdot a^0 b + 1 \cdot a^1 b^0\end{aligned}$$

Beweis dieser Formel nach der vollständigen Induktion:

Zuerst soll bewiesen werden, dass die Formel

$$(26) \quad \binom{n+1}{k} = \binom{n}{k-1} + \binom{n}{k}$$

tatsächlich richtig ist, also die Koeffizienten des Pascalschen Dreiecks liefert:

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

Also sind das Pascalsche Dreieck und seine Schreibweise in » $n$  über  $k$ « strukturgleich. Nun der Beweis der Formel. Behauptung: folgende Formel ist richtig:

$$\begin{aligned}(a+b)^n &= \binom{n}{n} a^n + \binom{n}{n-1} a^{n-1} b + \binom{n}{n-2} a^{n-2} b^2 + \dots + \binom{n}{1} a b^{n-1} + \binom{n}{0} b^n \\ &= \binom{n}{0} a^n + \binom{n}{1} a^{n-1} b + \binom{n}{2} a^{n-2} b^2 + \dots + \binom{n}{i} a^{n-i} b^i + \dots\end{aligned}$$

Sie sei richtig für  $k$ . Dann wird gezeigt nach der vollständigen Induktion, dass sie auch für  $k+1$  richtig ist. Wenn also

$$(a+b)^k = \binom{k}{0} a^k + \binom{k}{1} a^{k-1} b + \dots + \binom{k}{i} a^{k-i} b^i + \dots + \binom{k}{k} b^k$$

richtig ist, so muss gezeigt werden, dass auch folgende Formel richtig ist:

$$(a+b)^{k+1} = \binom{k+1}{0} a^{k+1} + \binom{k+1}{1} a^k b + \dots + \binom{k+1}{i} a^{k+1-i} b^i + \dots + \binom{k+1}{k+1} b^{k+1}$$

Um nun beide Formeln ineinander umzuformen, setze man die Richtigkeit der Formel für  $k$  voraus und multipliziere sie mit  $(a+b)$ , denn  $(a+b)^k \cdot (a+b) = (a+b)^{k+1}$ . Also:

$$\begin{aligned}
 (a+b)^k \cdot (a+b) &= (a+b)^{k+1} \\
 &= (a+b) \left\{ \binom{k}{0} a^k + \binom{k}{1} a^{k-1} b + \dots + \binom{k}{i} a^{k-i} b^i + \dots + \binom{k}{k} b^k \right\} \\
 &= a \cdot \left\{ \binom{k}{0} a^k + \binom{k}{1} a^{k-1} b + \dots + \binom{k}{i} a^{k-i} b^i + \dots + \binom{k}{k} b^k \right\} + \\
 &\quad b \cdot \left\{ \binom{k}{0} a^k + \binom{k}{1} a^{k-1} b + \dots + \binom{k}{i} a^{k-i} b^i + \dots + \binom{k}{k} b^k \right\} \\
 &= \binom{k+1}{0} a^{k+1} + \binom{k+1}{1} a^k b + \dots + \binom{k+1}{i} a^{k+1-i} b^i + \dots + \binom{k+1}{k+1} b^{k+1}
 \end{aligned}$$

Am Schluss wurde die Formel 26 in umgekehrter Richtung angewandt. Also wurde durch Ausmultiplizieren gezeigt, dass die allgemeine binomische Formel für  $k+1$  gilt, wenn sie für  $k$  gilt. Da sie außerdem für 1 gilt, ist sie nach dem Prinzip der vollständigen Induktion für jedes  $n \in \mathbb{N}$  richtig. Damit wurde die allgemeine binomische Formel bewiesen. Damit kann ohne weiteres gezeigt werden:

$$(1+x)^n \geq 1 + n \cdot x$$

denn

$$\begin{aligned}
 (1+x)^n &= \binom{n}{0} 1^n + \binom{n}{1} 1^{n-1} x + \binom{n}{2} 1^{n-2} x^2 + \dots + \binom{n}{n-1} 1 x^{n-1} + \binom{n}{n} x^n \\
 &\geq \binom{n}{0} 1^n + \binom{n}{1} 1^{n-1} x \\
 &\geq 1 + nx
 \end{aligned}$$

### 31. ÜBUNGSAUFGABEN

Entsprechend der bisherigen Kapitelgliederung werden hier Aufgaben von Prof. Eichner mit Lösung wiedergegeben. Die Reihenfolge der Aufgaben in einem Unterkapitel hat keine Bedeutung. Jede Aufgabe hat einen beschreibenden Titel und evtl. eine Angabe darüber, woher sie entnommen ist.

#### 31.1. Einige Bezeichnungen.

**31.2. Grundbegriffe der Mengenlehre.** Hinweis: Komplemente sind bezüglich der (i.d.R. angegebenen) Menge  $M$  zu bilden.

**31.2.1. Aufgabe 1.1.** Sei  $A = \{1, 2, 3, 4\}$ ,  $B = \{3, 5, 7\}$ ,  $C = \{6, 7\}$ ,  $M = \{1, 2, 3, 4, 5, 6, 7\}$ . Bilden sie  $A \cup B$ ,  $A \cap B$ ,  $A - B$ ,  $A \cap C$ ,  $\emptyset \cap A$ ,  $C \cap M$ ,  $\overline{A}$ ,  $\overline{B \cap C}$ .

Lösung.

- $A \cup B = \{1, 2, 3, 4, 5, 7\}$
- $A \cap B = \{3\}$
- $A \setminus B = \{1, 2, 4\}$
- $A \cap C = \emptyset$
- $\emptyset \cap A = \emptyset$
- $C \cap M = C$
- $\overline{A} = M \setminus A = \{5, 6, 7\}$
- $\overline{B \cap C} = \{1, 2, 4, 6\} \cap \{1, 2, 3, 4, 5\} = \{1, 2, 4\}$ . Auch zu lösen mit DeMorgan und Ausklammern:  
 $\overline{B \cap C} = (\overline{M \setminus B}) \cap (\overline{M \setminus C}) = M \setminus (B \cup C)$ . Weitere Alternative:  
 $\overline{B \cap C} = \overline{B \cup C} = M - (B \cup C) = \{1, 2, 4\}$

**31.2.2. Aufgabe 2.** Veranschaulichen sie jede Seite der folgenden Gleichungen in einem entsprechenden Venn-Diagramm.

- a):  $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$
- b):  $\overline{A \cup B} = \overline{A} \cap \overline{B}$
- c):  $\overline{A \cap B} = \overline{A} \cup \overline{B}$

**31.2.3. Aufgabe 3.** Seien  $A, B \subset M$ . Vereinfachen Sie:

- a):  $A \cap (A \cup B) = (A \cap A) \cup (A \cap B) = A \cup (A \cap B) = A$
- b):  $(A \cap B) \cup (\overline{A} \cap B) = (B \cap A) \cup (B \cap \overline{A}) = B \cap (A \cup \overline{A}) = B \cap M = B$
- c):  $A \cup \overline{B \cap A} = A \cup (\overline{B} \cup \overline{A}) = M$

#### 31.3. Aussagenlogik.

Rechenweg  
aufschreiben  
sche mittlere  
mit berücksichtigen  
Nachvollziehen  
Papula  
matik für  
(Hinweis durch  
Löffler).

31.3.1. *Aufgabe 1.* Prüfen sie für jede der folgenden Aussagen, ob sie wahr oder falsch ist, wobei bei den zusammengesetzten Aussagen die Tabellen nach Kapitel 3.1 zu verwenden sind.

- a):  $8 > 7 \Leftrightarrow w$   
 b):  $5 < 3 \Leftrightarrow f$   
 c):  $(4 < 5) \wedge (12 > 7) \Leftrightarrow w \wedge w \Leftrightarrow w$   
 d):  $(8 < 7) \Leftrightarrow \bar{f} \Leftrightarrow w$   
 e):  $\bar{5} < \bar{7} \Leftrightarrow f$   
 f): wenn  $4 > 3$  dann ist 5 eine Primzahl  $\Leftrightarrow w \rightarrow w \Leftrightarrow w$   
 g): wenn  $4 < 3$  dann 4 eine Primzahl  $\Leftrightarrow f \rightarrow f \Leftrightarrow w$   
 h): wenn  $4 > 3$  dann 4 eine Primzahl  $\Leftrightarrow w \rightarrow f \Leftrightarrow f$   
 i):  $(5 > 9) \leftrightarrow (3 > 4) \Leftrightarrow f \leftrightarrow f \Leftrightarrow w$   
 j):  $(3 < 4) \vee (3 = 4) \Leftrightarrow w \vee f \Leftrightarrow w$   
 k):  $3 \leq 4 \Leftrightarrow w$  (entspricht der Aussage in Teilaufgabe j)  
 l):  $(5 > 9) \oplus (3 > 4) \Leftrightarrow f \oplus f \Leftrightarrow f$   
 m):  $\bar{5} > \bar{9} \oplus (3 > 4) = w \oplus f = w$   
 n): der Esel ist ein Schaf genau dann, wenn das Pferd ein Vogel ist  $\Leftrightarrow f \leftrightarrow f \Leftrightarrow w$

31.3.2. *Aufgabe 2.* Zeigen Sie: (es muss jeweils gezeigt werden, dass der Werteverlauf der rechten und linken Formel identisch sind, so dass = erfüllt ist. Hier geht es im Gegensatz zu der Aufgabe in Kapitel 31.3.1 um aussagenlogische Variablen, nicht um einzelne Werte davon!).

- a):  $A \oplus B = \overline{A \leftrightarrow B}$ .

Werteverlauf der rechten Formel:

| A | B | $A \oplus B$ |
|---|---|--------------|
| f | f | f            |
| f | w | w            |
| w | f | w            |
| w | w | f            |

Werteverlauf der linken Formel, berechnet in mehreren Schritten:

Abbildung noch nicht von Zeichnung übertragen

Freiwillige vor!

ABBILDUNG 31. Venn-Diagramm zu  $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$

Abbildung noch nicht von Zeichnung übertragen

Freiwillige vor!

ABBILDUNG 32. Venn-Diagramm zu  $\overline{A \cup B} = \bar{A} \cap \bar{B}$

Abbildung noch nicht von Zeichnung übertragen

Freiwillige vor!

ABBILDUNG 33. Venn-Diagramm zu  $\overline{A \cap B} = \bar{A} \cup \bar{B}$



| $A$ | $B$ | $A \leftrightarrow B$ | $\overline{A \leftrightarrow B}$ |
|-----|-----|-----------------------|----------------------------------|
| $f$ | $f$ | $w$                   | $f$                              |
| $f$ | $w$ | $f$                   | $w$                              |
| $w$ | $f$ | $f$                   | $w$                              |
| $w$ | $w$ | $w$                   | $f$                              |

Da beide Werteverläufe identisch sind, gilt  $A \oplus B = \overline{A \leftrightarrow B}$ .

**b):** Die Aufgabenstellung enthält einen Schreibfehler; korrigiert:  $A \leftrightarrow B = \overline{(\overline{A} \wedge B) \vee (A \wedge \overline{B})}$ . Dies zeigt, dass die Äquivalenz durch UND, ODER, NICHT ausgedrückt werden kann.

31.3.3. *Aufgabe 3.* Bei welcher Formel handelt es sich um Tautologien? Wenn eine Formel immer wahr ist, handelt es sich um eine Tautologie. Es muss also wieder der Werteverlauf der Formel überprüft werden, wobei die Berechnung wieder in mehreren Schritten geschieht.

| $A$ | $B$ | $A \rightarrow B$ | $A \wedge (A \rightarrow B)$ | $\phi := (A \wedge (A \rightarrow B)) \rightarrow A$ |
|-----|-----|-------------------|------------------------------|------------------------------------------------------|
| $f$ | $f$ | $w$               | $f$                          | $w$                                                  |
| $f$ | $w$ | $w$               | $f$                          | $w$                                                  |
| $w$ | $f$ | $f$               | $f$                          | $w$                                                  |
| $w$ | $w$ | $w$               | $w$                          | $w$                                                  |

Also ist  $\phi$  eine Tautologie.

| $A$ | $B$ | $A \rightarrow B$ | $(A \rightarrow B) \wedge A := \alpha$ |
|-----|-----|-------------------|----------------------------------------|
| $f$ | $f$ | $w$               | $f$                                    |
| $f$ | $w$ | $w$               | $f$                                    |
| $w$ | $f$ | $f$               | $f$                                    |
| $w$ | $w$ | $w$               | $w$                                    |

**b):**

Also ist  $\alpha$  keine Tautologie.

| $B$ | $\overline{B}$ | $\overline{B} \rightarrow B$ | $(\overline{B} \rightarrow B) \rightarrow B := \beta$ |
|-----|----------------|------------------------------|-------------------------------------------------------|
| $w$ | $f$            | $w$                          | $w$                                                   |
| $f$ | $w$            | $f$                          | $w$                                                   |

**c):**

Also ist  $\beta$  eine Tautologie.

31.3.4. *Wahr oder falsch?* Quelle: [9]. Prüfen Sie, ob die folgende Aussage im Sinne der Aussagenlogik wahr oder falsch ist: »Entweder ist  $7 < 0$  oder  $4 < 8$ «.

$$(7 < 0) \oplus (4 < 8) = f \oplus w = w$$

### 31.4. Beweismethoden.

31.4.1. *Aufgabe 1.* Seien  $x, y \in \mathbb{N}$ . Beweisen sie:

**a):**  $xy$  ungerade  $\Leftrightarrow x$  ungerade  $\wedge y$  ungerade

Hier müssen wieder zwei Teilsätze bewiesen werden:

- $xy$  ungerade  $\Rightarrow x$  ungerade  $\wedge y$  ungerade

Nur mit indirektem Beweis möglich. Wir nehmen also an: » $x$  ungerade  $\wedge y$  ungerade« ist falsch, äquivalent » $x$  gerade  $\vee y$  gerade«. Es sind aufgrund des » $\vee$ « 3 Fälle zu unterscheiden, wobei alle drei zur Voraussetzung » $xy$  ungerade« im Widerspruch stehen müssen. Dies ist der Fall, also ist die Annahme falsch, also ist unter der Voraussetzung » $xy$  ungerade« die Behauptung » $x$  ungerade  $\wedge y$  ungerade« richtig.

(1)

$$\begin{aligned}
 & x \text{ gerade} \wedge y \text{ ungerade} \\
 \Rightarrow & x = 2k \wedge y = 2k' + 1 \quad (k, k' \in \mathbb{Z}) \\
 \Rightarrow & xy = 2k(2k' + 1) = 2(2kk' + k) \\
 \Rightarrow & xy = 2k'' \\
 \Rightarrow & xy \text{ gerade} \\
 \Rightarrow & \text{Widerspruch zur Voraussetzung}
 \end{aligned}$$

(2)

$$\begin{aligned}
 & x \text{ ungerade} \wedge y \text{ gerade} \\
 \Rightarrow & \text{(analog zum 1. Fall)} \\
 \Rightarrow & \text{Widerspruch zur Voraussetzung}
 \end{aligned}$$

(3)

$$\begin{aligned}
 & x \text{ gerade} \wedge y \text{ gerade} \\
 \Rightarrow & x = 2k \wedge y = 2k' \\
 \Rightarrow & xy = 2k \cdot 2k' \\
 \Rightarrow & 2(2kk') \\
 \Rightarrow & xy = 2k'' \\
 \Rightarrow & xy \text{ gerade} \\
 \Rightarrow & \text{Widerspruch zur Voraussetzung}
 \end{aligned}$$

- $xy$  ungerade  $\Leftarrow x$  ungerade  $\wedge y$  ungerade  
Beweisbar mit direktem Beweis:

$$\begin{aligned}
 & x \text{ ungerade} \wedge y \text{ ungerade} \\
 \Rightarrow & x = 2k + 1 \wedge y = 2k' + 1 \\
 \Rightarrow & xy = (2k + 1)(2k' + 1) = 4kk' + 2k + 2k' + 1 \\
 \Rightarrow & xy = 2(2kk' + k + k') + 1 \\
 \Rightarrow & xy = 2k'' + 1 \\
 \Rightarrow & xy \text{ ungerade}
 \end{aligned}$$

b):

$$(27) \quad x^3 \text{ gerade} \Leftrightarrow x \text{ gerade}$$

Es sind zwei Teilsätze zu beweisen:

- $x^3 \text{ gerade} \Rightarrow x \text{ gerade}$   
Es muss der indirekte Beweis gewählt werden, wie üblich wenn man von einem komplexeren auf einen einfacheren Ausdruck schließen will. Also Annahme: » $x$  ungerade«.

$$\begin{aligned}
 & x \text{ ungerade} \\
 \Rightarrow & x = 2k + 1 \\
 \Rightarrow & x^3 = (2k + 1)^3 = 8k^3 + 12k^2 + 6k + 1 \\
 \Rightarrow & x^3 = 2(4k^3 + 6k^2 + 3k) + 1 \\
 \Rightarrow & x^3 = 2k' + 1 \\
 \Rightarrow & x^3 \text{ ungerade} \\
 \Rightarrow & \text{Widerspruch zur Voraussetzung} \\
 \Rightarrow & x^3 \text{ gerade}
 \end{aligned}$$

- $x^3 \text{ gerade} \Leftarrow x \text{ gerade}$   
Direkter Beweis:

$$\begin{aligned}
 & \text{sei } x \text{ gerade} \\
 \Rightarrow & x = 2k \\
 \Rightarrow & x^3 = 8k^3 = 2(4k^3) \\
 \Rightarrow & x^3 = 2k' \\
 \Rightarrow & x^3 \text{ gerade}
 \end{aligned}$$

31.4.2. *Aufgabe 2.* Unter Verwendung von Definition 21 sind folgende Behauptungen zu beweisen: Sei  $n$  eine natürliche Zahl. Es gilt:

a):  $n$  durch 3 teilbar  $\Leftrightarrow n^2$  durch 9 teilbar

Der Beweis besteht wieder aus zwei Teilbeweisen:

- $n$  durch 3 teilbar  $\Rightarrow n^2$  durch 9 teilbar  
Direkter Beweis:

$$\begin{aligned}
 & \text{sei } n \text{ durch 3 teilbar} \\
 \Rightarrow & n = 3k \quad (k \in \mathbb{Z}) \\
 \Rightarrow & n^2 = 9k^2 \\
 \Rightarrow & n^2 \text{ durch 9 teilbar}
 \end{aligned}$$

- $n$  durch 3 teilbar  $\Leftrightarrow n^2$  durch 9 teilbar

Indirekter Beweis. Also Annahme: » $n$  durch 3 teilbar« ist falsch, d.h.  $n$  nicht durch 3 teilbar.

$$\begin{aligned} & n \text{ nicht durch 3 teilbar} \\ \Rightarrow & n = 3k + 1 \vee n = 3k + 2 \end{aligned}$$

Es sind also zwei Fälle zu unterscheiden, die beide einen Widerspruch zur Voraussetzung ergeben müssen:

(1)

$$\begin{aligned} & n = 3k + 1 \\ \Rightarrow & n^2 = 9k^2 + 6k + 1 = 3(3k^2 + 2k) + 1 \\ \Rightarrow & n^2 = 3k' + 1 \\ \Rightarrow & n^2 \text{ nicht durch 3 teilbar} \\ \Rightarrow & n^2 \text{ nicht durch 9 teilbar} \\ \Rightarrow & \text{Widerspruch zur Voraussetzung} \end{aligned}$$

(2)

$$\begin{aligned} & n = 3k + 2 \\ \Rightarrow & n^2 = 9k^2 + 12k + 4 = 3(3k^2 + 4k + 1) + 1 \\ \Rightarrow & n^2 = 3k' + 1 \\ \Rightarrow & n^2 \text{ nicht durch 3, aber auch nicht durch 9 teilbar} \\ \Rightarrow & \text{Widerspruch zur Voraussetzung} \end{aligned}$$

**b):**  $n^2$  dividiert durch 5 hat den Rest 1 oder 4  $\Leftrightarrow n$  nicht durch 5 teilbar

Der Beweis besteht aus zwei Teilbeweisen:

- $n^2$  dividiert durch 5 hat den Rest 1 oder 4  $\Rightarrow n$  nicht durch 5 teilbar  
Indirekter Beweis, also Annahme: » $n$  durch 5 teilbar«.

$$\begin{aligned} \Rightarrow & n = 5k \\ \Rightarrow & n^2 = 25k^2 = 5(5k^2) \\ \Rightarrow & n^2 = 5k' \\ \Rightarrow & n^2 \text{ durch 5 teilbar} \\ \Rightarrow & \text{Widerspruch zur Voraussetzung} \end{aligned}$$

- $n^2$  dividiert durch 5 hat den Rest 1 oder 4  $\Leftrightarrow n$  nicht durch 5 teilbar  
Direkter Beweis:

$$\begin{aligned} & n \text{ nicht durch 5 teilbar} \\ \Rightarrow & n = 5k + r; r = 1, 2, 3, 4 \end{aligned}$$

Es müssen 4 Fälle unterschieden werden:

(1) Fall:  $n = 5k + 1$

$$n^2 = 25k^2 + 10k + 1 = 5(5k^2 + 2k) + 1$$

(2) Fall:  $n = 5k + 2$

$$n^2 = 25k^2 + 20k + 4 = 5(5k^2 + 4k) + 4$$

(3) Fall:  $n = 5k + 3$

$$n^2 = 25k^2 + 30k + 9 = 5(5k^2 + 6k + 1) + 4$$

(4) Fall:  $n = 5k + 4$

$$n^2 = 25k^2 + 40k + 16 = 5(5k^2 + 8k + 3) + 1$$

31.5.1. *Aufgabe 1.* Prüfen sie für jede der folgenden Abbildungen, ob sie surjektiv, injektiv, bijektiv ist.

- a):  $f : \mathbb{Z} \mapsto \mathbb{Z}$  mit  $f(z) = z^2, z \in \mathbb{Z}$   
 $f$  nicht surjektiv, da z.B. 3 und  $-1$  keine Urbilder haben.  
 $f$  nicht injektiv, da z.B.  $f(-2) = f(+2) = 4$ .  
 $f$  nicht bijektiv, da nicht injektiv (und nicht surjektiv).
- b):  $g : \mathbb{N} \mapsto \mathbb{N}$  mit  $g(n) = 2n, n \in \mathbb{N}$   
 $g$  ist nicht surjektiv, da z.B. 9 kein Urbild hat.  
 $g$  ist injektiv ( $2n_1 \neq 2n_2$  für  $n_1 \neq n_2$ ).  
 $g$  ist nicht bijektiv, da sie nicht surjektiv ist.
- c):  $h : \mathbb{R} \mapsto \mathbb{R}$  mit  $h(r) = 2r, r \in \mathbb{R}$   
 $h$  ist surjektiv, denn jedes  $r \in \mathbb{R}$  ist durch 2 teilbar.  
 $h$  ist injektiv, denn  $2r_1 \neq 2r_2$  für  $r_1 \neq r_2$ .  
 $h$  ist bijektiv, da sie surjektiv und injektiv ist.
- d):  $j : \mathbb{N} \mapsto \{0, 1, 2\}$  mit  $j(n) = n \bmod 3, n \in \mathbb{N}$   
 $j$  ist surjektiv, da jede der Zahlen 0, 1, 2 als Rest vorkommen.  
 $j$  ist nicht injektiv, da z.B. 2 und 5 auf 2 abgebildet werden.  
 $j$  ist nicht bijektiv, da sie nicht injektiv ist.

31.5.2. *Aufgabe 2.*

- a): Sei  $A = \{a, b, c\}$  und  $B = \{0, 1\}$ . Geben sie sämtliche Abbildungen von  $A$  in  $B$  an. Dies ist in Form einer Tabelle möglich; es entsteht die Menge der dreistelligen Dualzahlen, wobei die Abbildungen entsprechend dem Wert der entstehenden Dualzahl bezeichnet wurden:

|       | $a$ | $b$ | $c$ |
|-------|-----|-----|-----|
| $f_0$ | 0   | 0   | 0   |
| $f_1$ | 0   | 0   | 1   |
| $f_2$ | 0   | 1   | 0   |
| $f_3$ | 0   | 1   | 1   |
| $f_4$ | 1   | 0   | 0   |
| $f_5$ | 1   | 0   | 1   |
| $f_6$ | 1   | 1   | 0   |
| $f_7$ | 1   | 1   | 1   |

- b): Sei  $A$  eine beliebige endliche Menge; sei  $n = |A|$ . Wieviele Abbildungen gibt es von  $A$  in die Menge  $\{0, 1\}$ ? Es gibt  $2^n$  Abbildungen  $f : A \mapsto B$ .

31.5.3. *Aufgabe 3.* Die folgende Aufgabe zeigt, wie Mengen und Teilmengen im Computer dargestellt werden können.

Sei  $A$  eine endliche Menge mit  $n$  Elementen. Sei  $T$  eine Teilmenge von  $A$ . Wir definieren eine Abbildung  $f_T : A \mapsto \{0, 1\}$  wie folgt:

$$f_T(a) = \begin{cases} 1 & \text{für } a \in T \\ 0 & \text{für } a \notin T \end{cases}$$

$f_T$  heißt charakteristische Abbildung von  $T$ . Legt man für die Elemente von  $A$  eine bestimmte Reihenfolge fest, dann wird mittels  $f_T$  jeder Teilmenge von  $A$  umkehrbar eindeutig eine  $n$ -stellige Dualzahl zugeordnet. Sei jetzt  $A = \{u, v, w, x, y, z\}$  mit der angegebenen Reihenfolge der Elemente.

- a): Welche Dualzahlen sind den Teilmengen  $T_1 = \{u, w, z\}$ ,  $T_2 = \{v\}$ ,  $T_3 = A$  zugeordnet? Zur Lösung verwendet man Wertetabellen, die den Wert der Abbildungen  $f_{T_i}$  für jedes  $a \in T_i$  angibt:

|           | $u$ | $v$ | $w$ | $x$ | $y$ | $z$ |
|-----------|-----|-----|-----|-----|-----|-----|
| $f_{T_1}$ | 1   | 0   | 1   | 0   | 0   | 1   |
| $f_{T_2}$ | 0   | 1   | 0   | 0   | 0   | 0   |
| $f_{T_3}$ | 1   | 1   | 1   | 1   | 1   | 1   |

- b): Welchen Teilmengen sind die Dualzahlen  $101101_2$ ,  $000001_2$ ,  $000000_2$  zugeordnet.

| $u$ | $v$ | $w$ | $x$ | $y$ | $z$ |                  |
|-----|-----|-----|-----|-----|-----|------------------|
| 1   | 0   | 1   | 1   | 0   | 1   | $\{u, w, x, z\}$ |
| 0   | 0   | 0   | 0   | 0   | 1   | $\{z\}$          |
| 0   | 0   | 0   | 0   | 0   | 0   | $\emptyset$      |

31.5.4. *Aufgabe 4.* Wieviele Teilmengen besitzt eine endliche Menge mit  $n$  Elementen? (Hinweis: vgl. Kapitel 31.5.3). Soviele, wie es  $n$ -stellige Dualzahlen gibt, nämlich  $2^n = 2^{|A|}$ . Deshalb heißt auch die Menge aller Teilmengen einer Menge ihre Potenzmenge.

31.5.5. *Aufgabe 5.* Wieviele Permutationen besitzt eine Menge mit 8 Elementen?  $8! = 40320$

### 31.6. Permutation.

### 31.7. Zahlen.

31.7.1. *Aufgabe 1.* Beweisen Sie an Hand der Grundgesetze der reellen Zahlen (7.1): Jede Gleichung  $a + x = b$  mit  $a, b \in \mathbb{R}$  besitzt genau eine Lösung in  $\mathbb{R}$ . Das erfordert zwei Teilbeweise:

- (1) »Jede Gleichung  $a + x = b$  mit  $a, b \in \mathbb{R}$  besitzt mindestens eine Lösung in  $\mathbb{R}$ «  
Wir zeigen:  $x = (-a) + b$  ist eine Lösung. Dazu einsetzen:

$$\begin{aligned} a + x &= b \\ \Rightarrow a + ((-a) + b) &= b \end{aligned}$$

Mit Assoziativgesetz gilt:

$$\Leftrightarrow (a + (-a)) + b = b$$

Mit Existenz des Inversen Elementes gilt:

$$\Leftrightarrow 0 + b = b$$

Mit Existenz des Neutralen Elementes gilt:

$$\Leftrightarrow b = b$$

- (2) »Jede Gleichung  $a + x = b$  mit  $a, b \in \mathbb{R}$  besitzt höchstens eine Lösung in  $\mathbb{R}$ «  
Dazu Annahme: Seien  $x_1, x_2$  Lösungen. Zu zeigen: dann muss gelten  $x_1 = x_2$ . Man geht von einer richtigen Aussage aus:

$$\begin{aligned} b &= b \\ \Leftrightarrow (-a) + b &= (-a) + b \end{aligned}$$

Mit  $a + x_1 = b \wedge a + x_2 = b$  gilt:

$$\Leftrightarrow (-a) + (a + x_1) = (-a) + (a + x_2)$$

Mit Assoziativgesetz gilt:

$$\Leftrightarrow ((-a) + a) + x_1 = ((-a) + a) + x_2$$

Mit Inversem Element gilt:

$$\Leftrightarrow 0 + x_1 = 0 + x_2$$

Mit Neutralem Element gilt:

$$\Leftrightarrow x_1 = x_2$$

qed.

31.7.2. *Aufgabe 2.* Machen Sie rational:

a):  $x = 32,1\bar{2}$

$$\begin{aligned} 100x &= 3212,\bar{2} \\ 10x &= 321,\bar{2} \\ \Rightarrow 90x &= 2891 \\ \Leftrightarrow x &= \frac{2891}{90} \end{aligned}$$

b):  $x = 0,56\overline{731}$

$$\begin{aligned} 100000x &= 56731,\overline{731} \\ 100x &= 56,\overline{731} \\ \Rightarrow 99900x &= 56675 \\ \Leftrightarrow x &= \frac{56675}{99900} = \frac{2267}{3996} \end{aligned}$$

31.7.3. *Aufgabe 3.* Zeigen Sie:  $0,42\bar{9} = 0,43$ .

$$\begin{aligned} 1000x &= 429,\bar{9} \\ 100x &= 42,\bar{9} \\ \Rightarrow 900x &= 387 \\ \Leftrightarrow x &= \frac{43}{100} = 0,43 \end{aligned}$$

31.7.4. *Aufgabe 4.* Beweisen Sie:  $x^3 = 2$  hat keine rationale Lösung.

Indirekter Beweis, also Annahme:  $x^3 = 2$  hat eine rationale Lösung, d.h.  $x = \frac{z}{n}$  sei eine Lösung. Als Voraussetzung, zu der wir einen Widerspruch erzeugen wollen, verwenden wir:  $x = \frac{z}{n}$  sei vollständig gekürzt.

$$\begin{aligned} x^3 &= 2 \\ \Rightarrow \frac{z^3}{n^3} &= 2 \\ \Rightarrow z^3 &= 2n^3 \\ \Rightarrow z^3 &\text{ ist eine gerade Zahl} \end{aligned}$$

Mit Gleichung 27 folgt:

$$(28) \quad \Rightarrow z \quad \text{gerade}$$

$$(29) \quad \Rightarrow z = 2k$$

Andererseits gilt:

$$\begin{aligned} \frac{z^3}{n^3} &= 2 \\ \Rightarrow n^3 &= \frac{z^3}{2} \end{aligned}$$

Mit Gleichung 29 folgt:

$$\begin{aligned} \Rightarrow \frac{8k^3}{2} &= 4k^3 = 2 \cdot (2k^3) \\ \Rightarrow n^3 &\text{ ist gerade} \end{aligned}$$

Mit Gleichung 27 folgt:

$$\begin{aligned} &\Rightarrow n \text{ ist gerade} \\ &\Rightarrow n = 2k' \\ &\Rightarrow z, n \text{ durch 2 teilbar} \\ &\Rightarrow \text{Widerspruch zur Voraussetzung, dass } \frac{z}{n} \text{ vollständig gekürzt sei} \\ &\Rightarrow \text{Annahme falsch, d.h. } x \text{ ist nicht rational} \end{aligned}$$

### 31.8. Mächtigkeit von Zahlenmengen.

### 31.9. Vollständige Induktion.

31.9.1. *Aufgabe 1.* Zeigen Sie durch vollständige Induktion:

a):

$$\sum_{k=1}^n \frac{1}{k(k+1)} = \frac{n}{n+1}, \quad n > 0$$

(1) Induktionsbeginn:  $n = 1$ . Die Behauptung gilt für  $n = 1$ :

$$\begin{aligned} \sum_{k=1}^1 \frac{1}{k(k+1)} &\stackrel{!}{=} \frac{1}{1+1} \\ \frac{1}{2} &= \frac{1}{2} \quad (w) \end{aligned}$$

(2) Induktionsannahme: Die Behauptung sei für ein beliebiges  $n \geq 1$  richtig. Diese Annahme gilt tatsächlich, nämlich für  $n = 1$ .

(3) Induktionsschluss von  $n$  auf  $n + 1$ : Zu zeigen ist unter Verwendung der Induktionsannahme:

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

Um Summenformeln mit vollständiger Induktion zu beweisen, spaltet man den Teil ab, der der Induktionsannahme entspricht, und wendet diese dann an:

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

qed.

**b):**

$$\sum_{k=1}^n k^2 = \frac{n(n+1)(2n+1)}{6}, \quad n > 0$$

(1) Induktionsbeginn:  $n = 1$ :

$$\begin{aligned}
 \sum_{k=1}^1 k^2 &\stackrel{!}{=} \frac{1(1+1)(2 \cdot 1 + 1)}{6} \\
 1 &= 1 \quad (w)
 \end{aligned}$$

(2) Induktionsannahme: Die Behauptung sei für ein  $n$  richtig.

(3) Induktionsschluss von  $n$  auf  $n+1$ : Zu zeigen ist:

$$\sum_{k=1}^{n+1} k^2 = \frac{(n+1)(n+2)(2n+3)}{6}$$

Diese Formel darf auch ausgerechnet werden, um schwierige Umformungen im Beweis zu vermeiden!

$$\begin{aligned}
 \sum_{k=1}^{n+1} k^2 &= \left( \sum_{k=1}^n k^2 \right) + (n+1)^2 \\
 &= \frac{n(n+1)(2n+1)}{6} + (n+1)^2 \\
 &= \frac{n(n+1)(2n+1) + 6(n+1)^2}{6} \\
 &= \frac{(n+1)(2n^2 + 7n + 6)}{6} \\
 &= \frac{(n+1)(n+2)(2n+3)}{6}
 \end{aligned}$$

31.9.2. *Aufgabe 2.* Zeigen Sie durch vollständige Induktion:

**a):**  $2^n > n$  für  $n \geq 1$

(1) Induktionsbeginn:  $n = 1$

$$\begin{aligned}
 2^1 &\stackrel{!}{>} 1 \\
 2 &> 1 \quad (w)
 \end{aligned}$$

(2) Induktionsannahme: Die Behauptung sei für ein  $n$  richtig.

(3) Induktionsschluss von  $n$  auf  $n+1$ : Zu zeigen ist unter Verwendung der Induktionsannahme:

$$2^{n+1} > n+1$$

Beweis:

$$2^{n+1} = 2^n \cdot 2$$

Unter Verwendung der Induktionsannahme und der Monotonie der Multiplikation gilt:

$$2^n \cdot 2 > n \cdot 2$$

Mit Monotonie der Addition ist  $n \geq 1 \Rightarrow n + n \geq n + 1$ :

$$\begin{aligned} 2^n \cdot 2 &> n \cdot 2 \geq n + 1 \\ \Rightarrow 2^n \cdot 2 &> n + 1 \end{aligned}$$

qed.

**b):**  $n! > 2^{n-1}$  für  $n \geq 3$

(1) Induktionsbeginn:  $n = 3$ :

$$\begin{aligned} 3! &\stackrel{!}{>} 2^{3-1} \\ 6 &> 4 \quad (w) \end{aligned}$$

(2) Induktionsannahme: Die Behauptung sei für ein  $n \geq 3$  richtig.

(3) Induktionsschluss von  $n$  auf  $n + 1$ : Zu zeigen ist unter Verwendung der Induktionsannahme:

$$(n + 1)! > 2^n$$

Beweis unter Verwendung der Induktionsannahme und Monotonie der Multiplikation:

$$\begin{aligned} (n + 1)! &= n! \cdot (n + 1) > 2^{n-1} \cdot (n + 1) > 2^{n-1} \cdot (1 + 1) = 2^n \\ \Rightarrow (n + 1)! &> 2^n \end{aligned}$$

qed.

**31.9.3. Aufgabe 3.** (Diese Aufgabe wird Ihnen in der Vorlesung »Datenstrukturen« möglicherweise noch einmal gestellt!) Ein ungerichteter Graph ist ein Paar  $G = (V, K)$ , wobei  $V$  eine endliche Menge ist und  $K \subset V \times V$ .  $V$  heißt Menge der Knoten,  $K$  Menge der ungerichteten Kanten. Für die Kantenmenge gilt: wenn  $(a, b) \in K$ , dann auch  $(b, a) \in K$ . Die Kanten  $(a, b)$  und  $(b, a)$  werden nicht unterschieden (genau diese Vereinbarung meint man, wenn man von ungerichteten Kanten spricht). Ferner vereinbaren wir, dass  $(a, a) \notin K$  für alle  $a \in V$ . Beweisen Sie durch vollständige Induktion: Die maximale Anzahl der Kanten in einem ungerichteten Graphen mit  $|V| = n$  Knoten ist  $|K| \leq \frac{n(n-1)}{2}$ .

(1) Induktionsbeginn:  $n = 1$ : Unter Verwendung von  $(a, a) \notin K$  ist:

$$\begin{aligned} |K| &\stackrel{!}{\leq} \frac{1 \cdot 0}{2} \\ 0 &= 0 \quad (w) \end{aligned}$$

(2) Induktionsannahme: Die Behauptung sei für ein  $n \geq 1$  richtig.

(3) Induktionsschluss von  $n$  auf  $n + 1$ : Sei  $V' = \{a_1, a_2, \dots, a_{n+1}\}$ . Zu zeigen ist unter Verwendung der Induktionsannahme:

$$|K'| \leq \frac{(n + 1)n}{2}$$

Beweis: Entferne einen Knoten  $x$  aus  $V'$

$$\begin{aligned} \Rightarrow V &= V' \setminus \{x\} \\ \Rightarrow K &= K' \setminus \{(a, b) \mid a = x \vee b = x\} \\ \Rightarrow |V| &= n \\ \Rightarrow |K| &\leq \frac{n(n-1)}{2} \end{aligned}$$

Bei Hinzufügen von  $x$  zu  $V$  mit  $|V| = n$  können maximal  $n$  Kanten  $(x, a)$  mit  $a \in V$  gebildet werden, nämlich mit den  $n$  bisherigen Knoten.

$$\begin{aligned} \Rightarrow |K'| &\leq \frac{n(n-1)}{2} + n \\ &\leq \frac{n(n+1)}{2} \end{aligned}$$

qed.

**31.9.4. Aufgabe 4.** Zeigen Sie durch vollständige Induktion:  $n$  verschiedene Geraden, die in einer Ebene durch einen gemeinsamen Punkt gehen, teilen die Ebene in  $2n$  Teile ( $n \geq 1$ ).



31.9.5. *Beweis von  $2^n > 2n + 1$ ,  $n \in \mathbb{N}$ ,  $n \geq n_0$ .* Quelle: [9]. Finden Sie das kleinste  $n_0$  und zeigen Sie durch vollständige Induktion:  $2^n > 2n + 1$ ,  $n \in \mathbb{N}$ ,  $n \geq n_0$ .

- (1) Induktionsbeginn. Für  $n_0 = 3$  ist  $2^3 > 2 \cdot 3 + 1 \Leftrightarrow 8 > 7$  erfüllt.
- (2) Induktionsannahme. Die Behauptung sei richtig für ein  $n \geq n_0$ .
- (3) Induktionsschluss von  $n$  auf  $n + 1$ . Zu zeigen:

$$\begin{aligned} A(n) &\stackrel{!}{\Rightarrow} A(n+1) \\ \Leftrightarrow 2^n > 2n + 1 &\stackrel{!}{\Rightarrow} 2^{n+1} > 2(n+1) + 1 \end{aligned}$$

Beweis mit direktem Beweis. Aufgrund der Induktionsannahme und der Monotonie der Multiplikation gilt:

$$\begin{aligned} 2^{n+1} = 2^n \cdot 2 &> (2n + 1) \cdot 2 \\ \Leftrightarrow 2^{n+1} &> 4n + 2 \\ \Leftrightarrow 2^{n+1} &> 2(n + 1) + 2n \end{aligned}$$

Mit  $2n > 1$  für  $n \in \mathbb{N}$ ,  $n \geq n_0$  folgt:

$$\Rightarrow 2^{n+1} > 2(n + 1) + 1$$

qed

### 31.10. Binominalkoeffizienten.

31.10.1. *Aufgabe 1.* Berechnen Sie manuell:

$$\begin{aligned} &\binom{9}{4} \\ &\binom{13}{10} \\ &\binom{100}{98} \\ &\binom{1000}{0} \\ &\binom{1000}{1000} \end{aligned}$$

31.10.2. *Aufgabe 2.* Wenden Sie den binomischen Satz auf  $(a - b)^7$  an.

31.10.3. *Aufgabe 3.* Berechnen Sie  $\frac{x^5 - 2^5}{x - 2}$ , wobei Sie sich auf eine Formel der Vorlesung stützen sollen, die Sie im Kapitel »Vollständige Induktion« kennengelernt haben.

### 31.11. Winkelfunktionen.

31.11.1. *Aufgabe 1.* Zeigen Sie an geeigneten Dreiecken:

$$\begin{aligned} \sin \frac{\pi}{4} &= \frac{\sqrt{2}}{2} \\ \sin \frac{\pi}{3} &= \frac{\sqrt{3}}{3} \\ \tan \frac{\pi}{6} &= \frac{\sqrt{3}}{3} \\ \tan \frac{\pi}{4} &= 1 \end{aligned}$$

31.11.2. *Aufgabe 2.* Berechnen Sie aus  $\cos 50^\circ = 0,6428$  die Werte von  $\sin x$ ,  $\cos x$ ,  $\tan x$ ,  $\cot x$  für  $x = 230^\circ$ . (Taschenrechner nur für Grundrechenarten und Radizieren).

31.11.3. *Aufgabe 3.* Drücken Sie in *rad* aus, und zwar als gekürzte rationale Vielfache von  $\pi$ :  $\alpha_1 = 10^\circ$ ,  $\alpha_2 = 315^\circ$ .

31.11.4. *Aufgabe 4.* Leiten Sie aus den Additionstheoremen (Gleichungen 12, 13) die Formeln her:

$$\begin{aligned} \tan(\alpha - \beta) &= \frac{\tan \alpha - \tan \beta}{1 + \tan \alpha \cdot \tan \beta} \\ \sin \gamma + \sin \delta &= 2 \cdot \sin \frac{\gamma + \delta}{2} \cdot \cos \frac{\gamma - \delta}{2} \end{aligned}$$

31.11.5. *Aufgabe 5.* Geben Sie die Lösungsmengen der Gleichungen an (als rationale Vielfache von  $\pi$ ):

a):  $\sin x = \frac{\sqrt{3}}{2}$ . Am Graphen der Sinusfunktion oder dem Einheitskreis erkennt man zwei Lösungen:  
 $\Rightarrow x = \frac{\pi}{3} \vee x = \pi - \frac{\pi}{3} = \frac{2\pi}{3}$ . Die Lösungsmenge ist mit der Periodizität  $2\pi$  also:

$$L = \left\{ \frac{\pi}{3} + k \cdot 2\pi \mid k \in \mathbb{Z} \right\} \cup \left\{ \frac{2\pi}{3} + k \cdot 2\pi \mid k \in \mathbb{Z} \right\}$$

b):  $\tan x = 1 \Rightarrow L = \left\{ \frac{\pi}{4} + k\pi \mid k \in \mathbb{Z} \right\}$

31.12. **Umkehrfunktion.**

31.13. **Die Arcusfunktionen.**

31.14. **Komplexe Zahlen.**

31.14.1. *Aufgabe 1.* Eine Zahl ist durch 4 teilbar, wenn ihre letzten beiden Ziffern durch 4 teilbar sind. Denn alle enthaltenen Potenzen  $10^2$  und höher sind ohnehin durch 4 teilbar.

a):  $i^{444} = 1$

b):  $i^{523} = i^{520} \cdot i^2 \cdot i = 1 \cdot -1 \cdot i = -i$

c):  $i^{2521} = i^{2520} \cdot i^1 = i$

d):  $i^{-1} = i^{-1} \cdot 1 = i^{-1} \cdot i^4 = i^3 = -i$

e):  $i^{-22} = i^{-22} \cdot 1 = i^{-22} \cdot i^{24} = i^2 = -1$

f):  $i^{-99} = i^{-99} \cdot 1 = i^{-99} \cdot i^{100} = i$

31.14.2. *Aufgabe 2.* Sei  $z_1 = 3 + 4i$ ,  $z_2 = -1 - 2i$ . Was ist:

a):  $z_1 + z_2 = 2 + 2i$

b):  $z_1 - z_2 = 4 + 6i$

c):  $|z_1| = \sqrt{3^2 + 4^2} = 5$

d):  $z_1^{-1} = \frac{1}{3+4i} = \frac{3-4i}{(3+4i)(3-4i)} = \frac{3-4i}{9+16} = \frac{3}{25} - \frac{4}{25}i$

e):  $z_1 \cdot z_2 = (3 + 4i)(-1 - 2i) = -3 - 6i - 4i + 8 = 5 - 10i$

f):  $\frac{z_1}{z_2} = \frac{(3+4i)(-1+2i)}{(-1-2i)(-1+2i)} = \frac{-3+6i-4i-8}{5} = -\frac{11}{5} + \frac{2}{5}i$

31.14.3. *Aufgabe 3.*

a):  $p(x) = x^4 + 6x^3 + 23x^2 + 34x + 26$ ,  $z_1 = -1 + i$

Die konjugiert komplexe Zahl zu  $z_1$  ist ebenfalls eine Lösung:  $z_2 = -1 - i$ . Man führt eine einzige Polynomdivision durch<sup>18</sup>

$$\begin{aligned} (x - z_1)(x - z_2) &= (x + 1 - i)(x + 1 + i) \\ &= ((x + 1) - i)((x + 1) + i) \\ &= (x + 1)^2 + 1 \\ &= x^2 + 2x + 2 \end{aligned}$$

durch:

$$(x^4 + 6x^3 + 23x^2 + 34x + 26) : (x^2 + 2x + 2) = x^2 + 4x + 13 := q_1(x)$$

Somit ist  $p(x) = (x - z_1)(x - z_2)q_1(x)$ . Lösen der quadratischen Gleichung  $q_1(x)$ :

$$\begin{aligned} q_1(x) &\stackrel{!}{=} 0 \\ x^2 + 4x + 13 &= 0 \\ \Rightarrow z_{3|4} &= -2 \pm \sqrt{4 - 13} \\ \Rightarrow z_3 = -2 + 3i \quad \vee \quad z_4 = -2 - 3i \end{aligned}$$

Damit ist die vollständige Linearfaktorzerlegung von  $p(x)$ :

$$\begin{aligned} p(x) &= (x - z_1)(x - z_2)(x - z_3)(x - z_4) \\ &= (x + 1 - i)(x + 1 + i)(x + 2 - 3i)(x + 2 + 3i) \end{aligned}$$

b): Was ist die Linearfaktorzerlegung des Polynoms:

$$h(x) = 2x^3 - 7x^2 + 8x - 3$$

Die erste Nullstelle  $z_1 = 1$  findet man durch Probieren. Man dividiert dann durch  $(x - 1)$  und erhält ein quadratisches Polynom:

$$(2x^3 - 7x^2 + 8x - 3) : (x - 1) = 2x^2 - 5x + 3 := q_1(x)$$

<sup>18</sup>es folgt ein gutes Verfahren zur Multiplikation, wenn zueinander konjugiert komplexe Zahlen vorkommen

Dieser quadratischen Gleichung spaltet man den höchsten Koeffizienten ab - er gehört mit zur Zerlegung von  $h(x)$  in Linearfaktoren:  $2(x^2 - \frac{5}{2}x + \frac{3}{2}) = 0$ . Nun löst man nach  $pq$ -Formel:

$$\begin{aligned} z_{2|3} &= \frac{5}{4} \pm \sqrt{\frac{25}{16} - \frac{3}{2}} = \frac{5}{4} \pm \sqrt{\frac{1}{16}} \\ &= \frac{5}{4} \pm \frac{1}{4} \\ \Rightarrow z_2 &= \frac{3}{2} \quad \vee \quad z_3 = 1 \end{aligned}$$

Damit ist die vollständige Linearfaktorzerlegung:

$$\begin{aligned} h(x) &= 2(x - z_1)(x - z_2)(x - z_3) \\ &= 2(x - 1)\left(x - \frac{3}{2}\right)(x - 1) \end{aligned}$$

31.14.4. *Aufgabe 4.* Setzen sie ohne Taschenrechner in trigonometrische Form um. Dazu stellt man sich die Zahl im Koordinatensystem vor.

**a):**  $z_1 = 5 = 5(\cos 0 + i \cdot \sin 0)$   
 $z_2 = -5 = 5(\cos \pi + i \cdot \sin \pi)$   
 $z_3 = 5i = 5(\cos \frac{\pi}{2} + i \cdot \sin \frac{\pi}{2})$   
 $z_4 = -5i = 5(\cos \frac{3\pi}{2} + i \cdot \sin \frac{3\pi}{2})$   
**b):**  $z_5 = 4 + 3i$ ,  $\tan \phi_0 = \frac{3}{4} \Rightarrow \arctan \phi_0 \approx 0,6435$ ,  $\Rightarrow z_5 = 5(\cos 0,6435 + i \cdot \sin 0,6435)$   
 $z_6 = -4 + 3i = 5(\cos(\pi - \phi_0) + i \cdot \sin(\pi - \phi_0)) = 5(\cos 2,4981 + i \cdot \sin 2,4981)$   
 $z_7 = -4 - 3i = 5(\cos(\pi + \phi_0) + i \cdot \sin(\pi + \phi_0)) = 5(\cos 3,7851 + i \cdot \sin 3,7851)$   
 $z_8 = 4 - 3i = 5(\cos(2\pi - \phi_0) + i \cdot \sin(2\pi - \phi_0)) = 5(\cos 5,6397 + i \cdot \sin 5,6397)$

31.14.5. *Aufgabe 5.* Rechnen Sie in kartesische Form um:

$$\begin{aligned} z &= 5\left(\cos \frac{\pi}{4} - i \cdot \sin \frac{\pi}{4}\right) \\ &= \frac{5\sqrt{2}}{2} - \frac{5\sqrt{2}}{2}i \end{aligned}$$

31.14.6. *Aufgabe 6.* Berechnen Sie in trigonometrischer Form:

$$\begin{aligned} (-1 - i)^{10} &= \left(\sqrt{2}\left(\cos \frac{5\pi}{4} + i \cdot \sin \frac{5\pi}{4}\right)\right)^{10} \\ &= \sqrt{2}^{10} \left(\cos\left(10 \frac{5\pi}{4}\right) + i \cdot \sin\left(10 \frac{5\pi}{4}\right)\right) \\ &= 32\left(\cos \frac{\pi}{2} + i \cdot \sin \frac{\pi}{2}\right) \\ &= 32i \end{aligned}$$

31.14.7. *Aufgabe 7.* Geben Sie alle Wurzeln von  $w^3 = \sqrt{3} - i$  in trigonometrischer Form.

$$\begin{aligned} \tan(2\pi - \phi_0) &= \frac{1}{\sqrt{3}} \Rightarrow \phi_0 = 2\pi - \frac{\pi}{6} = \frac{11\pi}{6} \\ \Rightarrow w^3 &= \sqrt{3} - i \\ &= (\sqrt{3} + 1) \left(\cos \frac{11\pi}{6} + i \cdot \sin \frac{11\pi}{6}\right) \\ \Rightarrow w_k &= \sqrt[3]{2} \left(\cos \frac{\frac{11\pi}{6} + k \cdot 2\pi}{3} + i \cdot \sin \frac{\frac{11\pi}{6} + k \cdot 2\pi}{3}\right) \end{aligned}$$

31.14.8. *kartesische Form in trigonometrische Form.* Quelle: [9]. Formen Sie die komplexe Zahl  $z = 16 + i \cdot 16\sqrt{3}$  in die trigonometrische Form um, wobei der Winkel als rationales Vielfaches von  $\pi$  anzugeben ist.

Lösungsverfahren siehe Kapitel 14.6.3. Gesucht ist eine Darstellungsform von  $z$ :

$$z = r(\cos \phi + i \cdot \sin \phi)$$

Wir erhalten:

$$\begin{aligned} r &= |z| = \sqrt{a^2 + b^2} = \sqrt{256 + 768} = 32 \\ \phi &= \arctan \frac{b}{a} = \arctan \sqrt{3} = \frac{\pi}{3} \\ \Rightarrow z &= 32 \left(\cos \left(\frac{\pi}{3} + 2k\pi\right) + i \cdot \sin \left(\frac{\pi}{3} + 2k\pi\right)\right), \quad k \in \mathbb{Z} \end{aligned}$$

31.14.9. *Division in kartesischer Form.* Quelle: [9]. Seien  $z_1 = 1 + i$ ,  $z_2 = 2 - 3i$ . Berechnen Sie  $\frac{z_1}{z_2}$ .

$$\frac{z_1}{z_2} = \frac{1+i}{2-3i} = \frac{(1+i)(2+3i)}{13} = \frac{-1+5i}{13} = -\frac{1}{13} + \frac{5}{13}i$$

31.14.10. *Eulersche Form in kartesische Form.* Quelle: [9]. Rechnen Sie die komplexe Zahl  $e^{i\frac{\pi}{4}}$  in die kartesische Form um.

$$e^{i\frac{\pi}{4}} = \cos \frac{\pi}{4} + i \sin \frac{\pi}{4} = \frac{1}{2}\sqrt{2} + \frac{1}{2}\sqrt{2}i$$

31.14.11. *Wurzeln in trigonometrischer Form.* Quelle: [10]. Geben Sie alle komplexen Lösungen der Gleichung  $x^3 = 27i$  in trigonometrischer Form an, wobei die Winkel als rationale Vielfache von  $\pi$  anzugeben sind.

$$\begin{aligned} x^3 &= 27i = 3^3 \left( \cos \frac{\pi}{2} + i \sin \frac{\pi}{2} \right) \\ \Rightarrow x_k &= 3 \left( \cos \frac{\frac{\pi}{2} + 2k\pi}{3} + i \sin \frac{\frac{\pi}{2} + 2k\pi}{3} \right) \\ \Rightarrow x_0 &= 3 \left( \cos \frac{\pi}{6} + i \sin \frac{\pi}{6} \right) \\ \vee x_1 &= 3 \left( \cos \frac{5\pi}{6} + i \sin \frac{5\pi}{6} \right) \\ \vee x_2 &= 3 \left( \cos \frac{9\pi}{6} + i \sin \frac{9\pi}{6} \right) = 3 \left( \cos \frac{3\pi}{2} + i \sin \frac{3\pi}{2} \right) \end{aligned}$$

### 31.15. Matrizen.

31.15.1. *Bewegungsmatrix für Drehung, Dehnung und Verschiebung.* Quelle: [9]. Ein Objekt ist um  $\frac{\pi}{4}$  um die  $x$ -Achse zu drehen, sodann um den Faktor 2 (Fixpunkt ist  $O$ ) zu dehnen und anschließend um den Vektor  $t = \begin{pmatrix} 1 & 0 & 1 \end{pmatrix}^T$  zu verschieben. Berechnen Sie die Bewegungsmatrix  $M$  in homogenen Koordinaten für die Gesamtbewegung.

Drehmatrix zur Drehung um die  $x$ -Achse:

$$M_1 = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & \cos \frac{\pi}{4} & -\sin \frac{\pi}{4} & 0 \\ 0 & \sin \frac{\pi}{4} & \cos \frac{\pi}{4} & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

Skalierungsmatrix zur Skalierung mit Faktor 2:

$$M_2 = \begin{pmatrix} 2 & 0 & 0 & 0 \\ 0 & 2 & 0 & 0 \\ 0 & 0 & 2 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

Translationsmatrix zur Verschiebung um den Vektor  $t = \begin{pmatrix} 1 & 0 & 1 \end{pmatrix}^T$ :

$$M_3 = \begin{pmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

Bewegungsmatrix der Gesamtbewegung:

$$\begin{aligned} M &= M_3 \cdot M_2 \cdot M_1 \\ &= \begin{pmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} 2 & 0 & 0 & 0 \\ 0 & 2 & 0 & 0 \\ 0 & 0 & 2 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & \cos \frac{\pi}{4} & -\sin \frac{\pi}{4} & 0 \\ 0 & \sin \frac{\pi}{4} & \cos \frac{\pi}{4} & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} 2 & 0 & 0 & 1 \\ 0 & 2 & 0 & 0 \\ 0 & 0 & 2 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & \frac{\sqrt{2}}{2} & -\frac{\sqrt{2}}{2} & 0 \\ 0 & \frac{\sqrt{2}}{2} & \frac{\sqrt{2}}{2} & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} 2 & 0 & 0 & 1 \\ 0 & \sqrt{2} & -\sqrt{2} & 0 \\ 0 & \sqrt{2} & \sqrt{2} & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix} \end{aligned}$$

31.15.2. *Bewegungsmatrix für Drehung um eine Gerade parallel zur x-Achse.* Quelle: [10]. Ein Objekt ist um  $\frac{\pi}{6}$  um eine Gerade  $g$  zu drehen, die parallel zur  $x$ -Achse durch den Punkt  $P_0 = (-1; -1; -1)$  verläuft. Berechnen Sie die Bewegungsmatrix  $M$  in homogenen Koordinaten.

Die Aufgabe wird hier als Beispiel für »Drehung um eine beliebige Gerade« gelöst. Sie kann wesentlich einfacher durch »Translation, Drehung um die  $x$ -Achse und Rücktranslation« gelöst werden. Geradengleichung mit normiertem Richtungsvektor:

$$g: \vec{X} = \vec{p}_0 + \lambda \vec{e} = \begin{pmatrix} -1 \\ -1 \\ -1 \\ 1 \end{pmatrix} + \lambda \begin{pmatrix} 1 \\ 0 \\ 0 \\ 1 \end{pmatrix}$$

Dann ist

$$M\left(\frac{\pi}{6}, g\right) = T(x_0, y_0, z_0) \cdot A\left(\frac{\pi}{6}\right) \cdot T(-x_0, -y_0, -z_0)$$

mit

$$T(x_0, y_0, z_0) = \begin{pmatrix} 1 & 0 & 0 & -1 \\ 0 & 1 & 0 & -1 \\ 0 & 0 & 1 & -1 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

$$T(-x_0, -y_0, -z_0) = \begin{pmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

$$A\left(\frac{\pi}{6}\right) = \begin{pmatrix} A^*\left(\frac{\pi}{6}\right) & & 0 \\ & & 0 \\ & & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

$$\begin{aligned} A^*\left(\frac{\pi}{6}\right) &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} + \left( \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} - \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \right) \cdot \left( \cos \frac{\pi}{6} + \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -1 \\ 0 & 1 & 0 \end{pmatrix} \sin \frac{\pi}{6} \right) \\ &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} + \begin{pmatrix} 0 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \cdot \left( \frac{1}{2}\sqrt{3} + \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -\frac{1}{2} \\ 0 & \frac{1}{2} & 0 \end{pmatrix} \right) \\ &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} + \begin{pmatrix} 0 & 0 & 0 \\ 0 & \frac{1}{2}\sqrt{3} & 0 \\ 0 & 0 & \frac{1}{2}\sqrt{3} \end{pmatrix} + \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -\frac{1}{2} \\ 0 & \frac{1}{2} & 0 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & \frac{1}{2}\sqrt{3} & -\frac{1}{2} \\ 0 & \frac{1}{2} & \frac{1}{2}\sqrt{3} \end{pmatrix} \end{aligned}$$

$$\Rightarrow A\left(\frac{\pi}{6}\right) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & \frac{1}{2}\sqrt{3} & -\frac{1}{2} & 0 \\ 0 & \frac{1}{2} & \frac{1}{2}\sqrt{3} & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

$$\begin{aligned} \Rightarrow M\left(\frac{\pi}{6}, g\right) &= T(x_0, y_0, z_0) \cdot A\left(\frac{\pi}{6}\right) \cdot T(-x_0, -y_0, -z_0) \\ &= \begin{pmatrix} 1 & 0 & 0 & -1 \\ 0 & 1 & 0 & -1 \\ 0 & 0 & 1 & -1 \\ 0 & 0 & 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & \frac{1}{2}\sqrt{3} & -\frac{1}{2} & 0 \\ 0 & \frac{1}{2} & \frac{1}{2}\sqrt{3} & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 & 0 & -1 \\ 0 & \frac{1}{2}\sqrt{3} & -\frac{1}{2} & -1 \\ 0 & \frac{1}{2} & \frac{1}{2}\sqrt{3} & -1 \\ 0 & 0 & 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & \frac{1}{2}\sqrt{3} & -\frac{1}{2} & (\frac{1}{2}\sqrt{3} - \frac{3}{2}) \\ 0 & \frac{1}{2} & \frac{1}{2}\sqrt{3} & (\frac{1}{2}\sqrt{3} - \frac{1}{2}) \\ 0 & 0 & 0 & 1 \end{pmatrix} \end{aligned}$$

### 31.16. Determinanten.

31.16.1. *Determinante einer  $4 \times 4$ -Matrix.* Quelle: [9]. Berechnen Sie die Determinante.

Entwicklung zuerst nach der dritten Spalte:

$$\begin{aligned} \begin{vmatrix} 3 & 2 & 0 & 4 \\ 1 & 4 & 0 & 0 \\ 111 & 701 & -1 & 1799 \\ 0 & -1 & 0 & 1 \end{vmatrix} &= -1 \begin{vmatrix} 3 & 2 & 4 \\ 1 & 4 & 0 \\ 0 & -1 & 1 \end{vmatrix} \\ &= -1 \cdot \left( 3 \cdot \begin{vmatrix} 4 & 0 \\ -1 & 1 \end{vmatrix} - 1 \cdot \begin{vmatrix} 2 & 4 \\ -1 & 1 \end{vmatrix} \right) \\ &= -1 \cdot (3 \cdot 4 - 1 \cdot 6) \\ &= -6 \end{aligned}$$

31.16.2. *Volumen eines Spats, Links- oder Rechtssystem?* Quelle: [9]. Gegeben sind die Vektoren  $\vec{a} = (1; -1; 2)^T$ ,  $\vec{b} = (1; 2; 0)^T$ ,  $\vec{c} = (3; -2; 1)^T$ .

a): Berechnen Sie das Volumen des von den Vektoren  $\vec{a}$ ,  $\vec{b}$ ,  $\vec{c}$  aufgespannten Spats.

$$\begin{aligned} V &= \left| (\vec{a}, \vec{b}, \vec{c}) \right| \\ &= \left| (\vec{a} \times \vec{b}) \cdot \vec{c} \right| \\ &= \left| \begin{pmatrix} -4 \\ 2 \\ 3 \end{pmatrix} \cdot \begin{pmatrix} 3 \\ -2 \\ 1 \end{pmatrix} \right| \\ &= |-12 - 4 + 3| \\ &= |-13| \\ &= 13 \end{aligned}$$

b): Bilden die Vektoren ein Links- oder Rechtssystem?

Da  $(\vec{a}, \vec{b}, \vec{c}) = -13$  negativ ist, bilden  $\vec{a}$ ,  $\vec{b}$ ,  $\vec{c}$  ein Linkssystem.

31.16.3. *Determinante einer  $5 \times 5$ -Matrix.* Quelle: [10]. Berechnen Sie die Determinante. Entwicklung zuerst nach der zweiten Zeile:

$$\begin{aligned} \begin{vmatrix} 0 & -17 & 0 & 1 & 0 \\ 0 & 4 & 0 & 0 & 0 \\ 3 & 329 & 3 & 7 & 0 \\ 1 & -35 & 0 & 4 & 4 \\ 1 & 66 & 1 & 76 & 1 \end{vmatrix} &= 4 \cdot \begin{vmatrix} 0 & 0 & 1 & 0 \\ 3 & 3 & 7 & 0 \\ 1 & 0 & 4 & 4 \\ 1 & 1 & 76 & 1 \end{vmatrix} \\ &= 4 \cdot \begin{vmatrix} 3 & 3 & 0 \\ 1 & 0 & 4 \\ 1 & 1 & 1 \end{vmatrix} \\ &= 4 \cdot \left( 3 \begin{vmatrix} 0 & 4 \\ 1 & 1 \end{vmatrix} - 3 \begin{vmatrix} 1 & 4 \\ 1 & 1 \end{vmatrix} \right) \\ &= 4 \cdot (3(-4) - 3(1-4)) \\ &= -12 \end{aligned}$$

### 31.17. Vektoren.

### 31.18. Skalarprodukt.

### 31.19. Vektorprodukt.

31.19.1. *Rechte oder linke Halbebene, Flächeninhalt.* Quelle: [10]. Sei  $\vec{a} = (1; -1000)^T$ ,  $\vec{b} = (-2; 1000)^T$ .

- Berechnen Sie, ob  $\vec{b}$  in die linke oder rechte Halbebene von  $\vec{a}$  zeigt.

$$D = \begin{vmatrix} a_1 & a_2 \\ b_1 & b_2 \end{vmatrix} = \begin{vmatrix} 1 & -1000 \\ -2 & 1000 \end{vmatrix} = -1000 < 0$$

Da  $D < 0$ , zeigt  $\vec{b}$  in die rechte Halbebene von  $\vec{a}$ .

- Geben Sie den Flächeninhalt des von  $\vec{a}$  und  $\vec{b}$  aufgespannten Parallelogramms an.

$$A = |D| = 1000$$

### 31.20. Spatprodukt.

31.20.1. *Rechts- oder Linkssystem, Spatvolumen.* Quelle: [10]. Gegeben sind die Vektoren  $\vec{a} = (1; 1; 1)^T$ ,  $\vec{b} = (0; 2; 1)^T$ ,  $\vec{c} = (-1; 0; -3)^T$ . Berechnen Sie, ob die Vektoren in der Reihenfolge  $\vec{a}$ ,  $\vec{b}$ ,  $\vec{c}$  ein Rechts- oder Linkssystem bilden.

$$\begin{aligned} [\vec{a}, \vec{b}, \vec{c}] &= \begin{vmatrix} 1 & 1 & 1 \\ 0 & 2 & 1 \\ -1 & 0 & -3 \end{vmatrix} = \begin{vmatrix} 2 & 1 \\ 0 & -3 \end{vmatrix} - \begin{vmatrix} 0 & 1 \\ -1 & -3 \end{vmatrix} + \begin{vmatrix} 0 & 2 \\ -1 & 0 \end{vmatrix} \\ &= -6 - 1 + 2 \\ &= -5 \end{aligned}$$

Da  $[\vec{a}, \vec{b}, \vec{c}]$  negativ ist, bilden  $\vec{a}$ ,  $\vec{b}$ ,  $\vec{c}$  in dieser Reihenfolge ein Linkssystem.

Geben Sie das Volumen des von  $\vec{a}$ ,  $\vec{b}$ ,  $\vec{c}$  aufgespannten Spats an.

$$V = |[\vec{a}, \vec{b}, \vec{c}]| = |-5| = 5$$

### 31.21. Geraden.

31.21.1. *Windschiefheit und Abstand von Geraden.* Quelle: [9]. Gegeben die Geraden

$$\begin{aligned} g_1 : \vec{X} &= \vec{p}_1 + \lambda \vec{a}_1 \\ &= (1; 1; 2)^T + \lambda (-1; 2; 1)^T \end{aligned}$$

$$\begin{aligned} g_2 : \vec{X} &= \vec{p}_2 + \mu \vec{a}_2 \\ &= (0; 1; 2)^T + \mu (1; 1; 1)^T \end{aligned}$$

**a):** Zeigen Sie: Die Geraden sind windschief.

$g_1$ ,  $g_2$  sind windschief, wenn  $\vec{a}_1$ ,  $\vec{a}_2$ ,  $\vec{p}_1 - \vec{p}_2$  linear unabhängig sind, d.h. wenn  $|(a_1, a_2, p_1 - p_2)| \neq 0$ .

$$\begin{vmatrix} -1 & 1 & 1 \\ 2 & 1 & 0 \\ 1 & 1 & 0 \end{vmatrix} = \begin{vmatrix} 2 & 1 \\ 1 & 1 \end{vmatrix} = 2 - 1 = 1$$

qed.  $(n, n)$ -Determinanten stellen ja Volumina im  $n$ -dimensionalen Vektorraum dar. Lineare Unabhängigkeit von drei Vektoren im  $\mathbb{R}^3$  bedeutet aber, dass sie nicht komplanar sind, d.h. dass das Volumen des von ihnen aufgespannten Spats nicht 0 ist.

**b):** Berechnen Sie den Abstand der Geraden. Nach Gleichung 21 ist der Abstand windschiefer Geraden:

$$\begin{aligned} d(g_1, g_2) &= \frac{|[(\vec{p}_2 - \vec{p}_1), \vec{a}_1, \vec{a}_2]|}{|\vec{a}_1 \times \vec{a}_2|} \\ &= \frac{\left| \left[ \begin{pmatrix} -1 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} -1 \\ 2 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \right] \right|}{\left| \begin{pmatrix} -1 \\ 2 \\ 1 \end{pmatrix} \times \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \right|} \\ &= \frac{\left\| \begin{vmatrix} -1 & -1 & 1 \\ 0 & 2 & 1 \\ 0 & 1 & 1 \end{vmatrix} \right\|}{\left\| \begin{vmatrix} 1 \\ 2 \\ -3 \end{vmatrix} \right\|} \\ &= \frac{|-1|}{\sqrt{14}} = \frac{1}{\sqrt{14}} = 14^{-\frac{1}{2}} \end{aligned}$$

31.21.2. *Schnittpunkt und Schnittwinkel von Geraden.* Quelle: [10]. Gegeben die Geraden

$$g_1 : \vec{X} = \vec{p}_1 + \lambda \vec{a}_1 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} -2 \\ -2 \\ -2 \end{pmatrix}$$

$$g_2 : \vec{X} = \vec{p}_2 + \mu \vec{a}_2 = \begin{pmatrix} 3 \\ -1 \\ 2 \end{pmatrix} + \mu \begin{pmatrix} -1 \\ 4 \\ 0 \end{pmatrix}$$

Berechnen Sie den Schnittpunkt. Dazu:

$$\begin{aligned} g_1 &= g_2 \\ \Rightarrow \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} -2 \\ -2 \\ -2 \end{pmatrix} &= \begin{pmatrix} 3 \\ -1 \\ 2 \end{pmatrix} + \mu \begin{pmatrix} -1 \\ 4 \\ 0 \end{pmatrix} \\ \Leftrightarrow \mu \begin{pmatrix} 1 \\ -4 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} -2 \\ -2 \\ -2 \end{pmatrix} &= \begin{pmatrix} 3 \\ -2 \\ 2 \end{pmatrix} \\ \Rightarrow \lambda_s &= -1 \end{aligned}$$

Normalerweise berechnet man nun noch  $\mu_s$  und setzt  $\lambda_s$  und  $\mu_s$  in die noch unbenutzte Gleichung zur Probe ein. Wir sparen uns das hier, weil die Aufgabenstellung impliziert, dass ein Schnittpunkt existiert. Setze  $\lambda_s = -1$  in  $g_1$  ein, um den Schnittpunkt zu erhalten:

$$\begin{aligned} \vec{s} &= \vec{p}_1 + \lambda_s \vec{a}_1 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} - \begin{pmatrix} -2 \\ -2 \\ -2 \end{pmatrix} = \begin{pmatrix} 2 \\ 3 \\ 2 \end{pmatrix} \\ \Rightarrow S(2; 3; 2) \end{aligned}$$

Berechnen Sie den Cosinus des Schnittwinkels  $\alpha$ .

$$\begin{aligned} \cos \phi &= \left| \frac{\vec{a}_1 \vec{a}_2}{a_1 a_2} \right| = \left| \frac{\begin{pmatrix} -2 \\ -2 \\ -2 \end{pmatrix} \begin{pmatrix} -1 \\ 4 \\ 0 \end{pmatrix}}{\sqrt{12}\sqrt{17}} \right| \\ &= \left| \frac{-6}{\sqrt{204}} \right| = \sqrt{\frac{9}{51}} = \frac{3}{\sqrt{51}} \end{aligned}$$

Berechnen Sie eine Gerade  $g : \vec{X} = \vec{p}_3 + \nu \vec{a}_3$ , die zu  $g_1$  und  $g_2$  senkrecht steht und durch den Schnittpunkt geht.

$$\begin{aligned} \vec{p}_3 &= \vec{s} \\ \vec{a}_3 \perp \vec{a}_1 &\wedge \vec{a}_3 \perp \vec{a}_2 \\ \Leftrightarrow \vec{a}_3 \vec{a}_1 &= 0 \wedge \vec{a}_3 \vec{a}_2 = 0 \\ \Leftrightarrow -2a_{31} - 2a_{32} - 2a_{33} &= 0 \wedge -a_{31} + 4a_{32} = 0 \\ \Rightarrow a_{32} &= t \\ \wedge a_{31} &= 4t \\ \wedge a_{33} &= -5t \end{aligned}$$

Also ist  $\vec{a}_3$  mit einem gewählten  $t = 1$ :

$$\begin{aligned} \vec{a}_3 &= \begin{pmatrix} 4 \\ 1 \\ -5 \end{pmatrix} \\ \Rightarrow g : \vec{X} &= \begin{pmatrix} 2 \\ 3 \\ 2 \end{pmatrix} + \nu \begin{pmatrix} 4 \\ 1 \\ -5 \end{pmatrix} \end{aligned}$$

31.22. **Ebenen.**



31.22.1. Gerade senkrecht zu einer Ebene, Ebene in Koordinatenform, Halbräume einer Ebene. Quelle: [10]. Sei

$$E : \vec{X} = \vec{p} + \lambda \vec{a} + \mu \vec{b} = \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix} + \lambda \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} + \mu \begin{pmatrix} -1 \\ 0 \\ 2 \end{pmatrix}$$

- Berechnen Sie eine Gerade  $g : \vec{X} = \vec{q} + \nu \vec{n}$  in Parameterform, die durch den Punkt  $R = (4; 1; 0)$  geht und senkrecht zu  $E$  verläuft.

$$\vec{q} = \vec{r} = \begin{pmatrix} 4 \\ 1 \\ 0 \end{pmatrix}$$

$\vec{n} \perp \vec{a} \wedge \vec{n} \perp \vec{b}$ , wir erhalten  $\vec{n}$  also einfach über das Vektorprodukt:

$$\vec{n} = \vec{a} \times \vec{b} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \times \begin{pmatrix} -1 \\ 0 \\ 2 \end{pmatrix} = \begin{pmatrix} 2 \\ -3 \\ 1 \end{pmatrix}$$

$$\Rightarrow g : \vec{X} = \vec{q} + \nu \vec{n} = \begin{pmatrix} 4 \\ 1 \\ 0 \end{pmatrix} + \nu \begin{pmatrix} 2 \\ -3 \\ 1 \end{pmatrix}$$

- Geben Sie die Koordinatenform von  $E$  an. Über Normalenform:

$$\begin{aligned} E : \vec{n}(\vec{x} - \vec{p}) &= 0 \\ \Rightarrow E : \begin{pmatrix} 2 \\ -3 \\ 1 \end{pmatrix} \left( \vec{x} - \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix} \right) &= 0 \\ \Leftrightarrow E : 2(x_1 - 1) - 3x_2 + x_3 + 1 &= 0 \\ \Leftrightarrow 2x_1 - 3x_2 + x_3 - 1 &= 0 \end{aligned}$$

Man kann auch direkt von der Gleichung  $\vec{n}\vec{x} - \vec{n}\vec{p} = 0$  ausgehen und den Rechenweg so etwas kürzen.

- Berechnen Sie, ob  $R$  und der Koordinatenursprung  $O$  auf der gleichen Seite von  $E$  liegen. Dies ist der Fall, wenn  $\vec{n}$  zu  $R$  und  $O$  gleich orientiert ist. Dazu in Normalenform von  $E$  einsetzen:

$$\vec{n}(\vec{o} - \vec{p}) = \begin{pmatrix} 2 \\ -3 \\ 1 \end{pmatrix} \begin{pmatrix} -1 \\ 0 \\ 1 \end{pmatrix} = -1 < 0$$

$$\begin{aligned} \vec{n}(\vec{r} - \vec{p}) &= \begin{pmatrix} 2 \\ -3 \\ 1 \end{pmatrix} \left( \begin{pmatrix} 4 \\ 1 \\ 0 \end{pmatrix} - \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix} \right) \\ &= \begin{pmatrix} 2 \\ -3 \\ 1 \end{pmatrix} \begin{pmatrix} 3 \\ 1 \\ 1 \end{pmatrix} = 6 - 3 + 1 = 4 > 0 \end{aligned}$$

Also liegen  $R$  und  $O$  nicht im selben Halbraum von  $E$ .

31.22.2. Gerade senkrecht zu einer Ebene, Ebene in Koordinatenform. Quelle: [9]. Gegeben die Ebene  $E : \vec{X} = \vec{p} + \lambda \vec{a} + \mu \vec{b} = (2; 1; -5)^T + \lambda (2; 0; 1)^T + \mu (1; 1; 1)^T$ .

- a):** Berechnen Sie eine Gerade  $g$  in Parameterform, die durch  $Q = (-1; -1; 2)$  geht und senkrecht durch  $E$  verläuft.

Gesucht ist  $g : \vec{X} = \vec{q} + \nu \vec{c}$  mit  $\vec{q} = (-1; -1; 2)^T$  und  $\vec{c} \perp \vec{a} \wedge \vec{c} \perp \vec{b}$ . Ermittle  $\vec{c}$ :

$$\begin{aligned} \begin{pmatrix} a_1 & a_2 & a_3 \\ b_1 & b_2 & b_3 \end{pmatrix} \vec{c} &= \begin{pmatrix} 0 \\ 0 \end{pmatrix} \\ \Rightarrow c_3 &= t \\ \wedge c_1 &= -\frac{1}{2}t \\ \wedge c_2 &= -\frac{1}{2}t \end{aligned}$$

$$\Rightarrow g : \vec{X} = \begin{pmatrix} -1 \\ -1 \\ 2 \end{pmatrix} + \nu \begin{pmatrix} -1 \\ -1 \\ 2 \end{pmatrix}$$

**b):** Geben Sie die Koordinatenform von  $E$  an.

Die Normalenform von  $E$  ist, unter Verwendung des in Aufgabenteil a) berechneten Normalenvektors  $\vec{c}$ :

$$\begin{aligned} E : \vec{c}(\vec{x} - \vec{p}) &= 0 \\ \Rightarrow E : \begin{pmatrix} -1 \\ -1 \\ 2 \end{pmatrix} \left( \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} - \begin{pmatrix} 2 \\ 1 \\ -5 \end{pmatrix} \right) &= 0 \end{aligned}$$

Ermittlung der Koordinatenform durch einfaches Ausrechnen:

$$\begin{aligned} E : -(x_1 - 2) - (x_2 - 1) + 2(x_3 + 5) &= 0 \\ \Leftrightarrow E : -x_1 + 2 - x_2 + 1 + 2x_3 + 10 &= 0 \\ \Leftrightarrow E : -x_1 - x_2 + 2x_3 + 13 &= 0 \end{aligned}$$

31.22.3. *Ebene in Hesseform, Abstand Punkt-Ebene, Halbräume einer Ebene.* Quelle: [9]. Sei  $E : 2x_1 - 2x_2 + 1x_3 = 6$ .

**a):** Geben Sie  $E$  in Hesseform an.

$$\begin{aligned} \vec{n}^0 &= \frac{\begin{pmatrix} 2 \\ -2 \\ 1 \end{pmatrix}}{\sqrt{9}} \\ E : \begin{pmatrix} \frac{2}{3} \\ -\frac{2}{3} \\ \frac{1}{3} \end{pmatrix} \left( \vec{x} - \begin{pmatrix} 0 \\ 0 \\ 6 \end{pmatrix} \right) &= 0 \end{aligned}$$

**b):** Geben Sie den Abstand des Ursprungs  $O$  von  $E$  an.

$$d(O, E) = \left| \vec{n}^0 \left( \vec{o} - \begin{pmatrix} 0 \\ 0 \\ 6 \end{pmatrix} \right) \right| = |-2| = 2$$

**c):** Berechnen Sie, ob  $Q = (-2; 0; 1)$  und  $O$  auf verschiedenen Seiten von  $E$  liegen.

$$d(Q, E) = \left| \vec{n}^0 \left( \vec{q} - \begin{pmatrix} 0 \\ 0 \\ 6 \end{pmatrix} \right) \right| = \left| -\frac{4}{3} - \frac{5}{3} \right| = |-3|$$

Da vor Ermittlung des Betrages in  $d(O, E)$  und  $d(Q, E)$  gleiches Vorzeichen bestand, ist  $\vec{n}^0$  zu  $Q$  und  $O$  gleich orientiert, d.h.  $Q$  und  $O$  liegen auf der gleichen Seite der Ebene  $E$ .

31.22.4. *Schnittgerade von Ebenen in Parameterform.* Quelle: [9]. Gegeben sind die Ebenen

$$\begin{aligned} E_1 : \vec{X} &= \vec{p}_1 + \lambda_1 \vec{a}_1 + \mu_1 \vec{b}_1 = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} + \lambda_1 \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} + \mu_1 \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} \\ (30) \quad E_2 : \vec{X} &= \vec{p}_2 + \lambda_2 \vec{a}_2 + \mu_2 \vec{b}_2 = \begin{pmatrix} -1 \\ -1 \\ 0 \end{pmatrix} + \lambda_2 \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} + \mu_2 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} \end{aligned}$$

Berechnen Sie die Schnittgerade in Parameterform. Setze  $E_1 = E_2$  und berechne einen Parameter in Abhängigkeit des anderen; so entsteht die Gleichung der Schnittgeraden.

|             |         |             |         |    |
|-------------|---------|-------------|---------|----|
| $\lambda_1$ | $\mu_1$ | $\lambda_2$ | $\mu_2$ |    |
| 1           | 0       | 0           | -1      | -2 |
| 0           | 1       | 0           | 0       | -1 |
| 1           | 0       | -1          | 0       | -1 |
| 1           | 0       | 0           | -1      | -2 |
| 0           | 1       | 0           | 0       | -1 |
| 0           | 0       | 1           | -1      | -1 |

$\Rightarrow \lambda_2 = -1 + \mu_2$  einsetzen in 30:

$$\begin{aligned} g : \vec{X} &= \begin{pmatrix} -1 \\ -1 \\ 0 \end{pmatrix} + (-1 + \mu_2) \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} + \mu_2 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} \\ &= \begin{pmatrix} -1 \\ -1 \\ 0 \end{pmatrix} + \begin{pmatrix} 0 \\ 0 \\ -1 \end{pmatrix} + \mu_2 \left( \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} + \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} \right) \\ &= \begin{pmatrix} -1 \\ -1 \\ -1 \end{pmatrix} + \nu \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} \end{aligned}$$

31.22.5. *Schnittpunkt von Gerade und Ebene.* Quelle: [10]. Berechnen Sie den Schnittpunkt der Geraden  $g$  mit der Ebenen  $E$ :

$$g : \vec{X} = \vec{q} + \lambda \vec{a} = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} + \lambda \begin{pmatrix} 2 \\ 1 \\ -3 \end{pmatrix}$$

$$\begin{aligned} E : n_1 x_1 + n_2 x_2 + n_3 x_3 &= \vec{n} \vec{p} \\ \Rightarrow E : 3x_1 + 2x_2 - x_3 &= 10 \end{aligned}$$

Setze die Teilgleichungen für die Koordinaten von Punkten auf  $g$  in  $E$  ein:

$$\begin{aligned} n_1 (q_1 + \lambda a_1) + n_2 (q_2 + \lambda a_2) + n_3 (q_3 + \lambda a_3) &= \vec{n} \vec{p} \\ \Rightarrow 3(2\lambda) + 2(\lambda) - 1(1 - 3\lambda) &= 10 \\ \Leftrightarrow 11\lambda &= 11 \\ \Leftrightarrow \lambda &= 1 \end{aligned}$$

Setze das erhaltene  $\lambda$  in  $g$  ein:

$$\begin{aligned} \vec{s} &= \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} + 1 \begin{pmatrix} 2 \\ 1 \\ -3 \end{pmatrix} = \begin{pmatrix} 2 \\ 1 \\ -2 \end{pmatrix} \\ \Rightarrow S &= (2; 1; -2) \end{aligned}$$

31.22.6. *Ebene in Hesseform, parallele Ebene erstellen, Abstand paralleler Ebenen.* Quelle: [10]. Sei  $E : 2x_1 - 3x_2 + 6x_3 = -21$ .

- Geben Sie  $E$  in Hesseform an.

$$\begin{aligned} E : \frac{2x_1 - 3x_2 + 6x_3}{-\sqrt{49}} &= \frac{-21}{-\sqrt{49}} \\ \Leftrightarrow -\frac{2}{7}x_1 + \frac{3}{7}x_2 - \frac{6}{7}x_3 &= 3 \end{aligned}$$

- Geben Sie eine zu  $E$  parallele Ebene  $E'$  (in Hesseform) an, die durch den Punkt  $Q = (1; 1; 1)$  geht.

$$\begin{aligned} E : \vec{n}^0 (\vec{x} - \vec{q}) &= 0 \\ \Leftrightarrow E' : \vec{n}^0 \vec{x} &= \vec{n}^0 \vec{q} \\ \Leftrightarrow -\frac{2}{7}x_1 + \frac{3}{7}x_2 - \frac{6}{7}x_3 &= -\frac{5}{7} \end{aligned}$$

- Welchen Abstand haben  $E$  und  $E'$ ?

$$d(E, E') = d(E, Q) = |\vec{n}^0 (\vec{q} - \vec{p})| = |\vec{n}^0 \vec{q} - \vec{n}^0 \vec{p}| = \left| -\frac{5}{7} - 3 \right| = \frac{26}{7}$$

31.23. **Spiegelungen.**

31.24. **Kreise und Kugeln.**

31.25. **Folgen und Reihen.**

31.25.1. *Ungleichung mit  $K$  und  $n_0$ .* Quelle: [9]. Berechnen Sie ein  $K$  und ein  $n_0$  mit  $\left| \frac{n^5}{3} - 2n^2 - \frac{1}{3} \right| \leq Kn^5$  für alle  $n \geq n_0$ . Da es hier um eine Folge  $(a_n)$  geht, ist sicher  $n_0 \geq 1 \Rightarrow n^5 > 0$ . Wir wählen in diesem Bereich  $n_0$  frei zu  $n_0 = 1$ .

$$\begin{aligned} \left| \frac{n^5}{3} - 2n^2 - \frac{1}{3} \right| &\stackrel{!}{\leq} Kn^5 \\ \Rightarrow n^5 \left| \frac{1}{3} - \frac{2}{n^3} - \frac{1}{3n^5} \right| &\stackrel{!}{\leq} Kn^5 \\ \Leftrightarrow \left| \frac{1}{3} - \frac{2}{n^3} - \frac{1}{3n^5} \right| &\leq \left( \frac{1}{3} + \frac{2}{n^3} + \frac{1}{3n^5} \right) \leq \left( \frac{8}{3} \right) \stackrel{!}{\leq} K \end{aligned}$$

$\frac{8}{3} \stackrel{!}{\leq} K$  ist für  $K = \frac{8}{3}$  erfüllt. Die Lösung ist also  $K = \frac{8}{3}$  für  $n \geq n_0 = 1$ . Bei einer solchen Aufgabe ist das Maximum einer Folge zu bestimmen, nicht ihr Wert für  $n \rightarrow \infty$ . Bei nicht monoton wachsenden Folgen wie dieser kann der Grenzwert für  $n \rightarrow \infty$  ja kleiner als das Maximum sein! Man wählt zuerst ein  $n_0$  und bestimmt das Maximum der Folge für dieses  $n_0$  dann, indem man schrittweise einfachere Ausdrücke findet, deren Wert die Folge nicht überschreiten kann:

- $\left| \frac{1}{3} - \frac{2}{n^3} - \frac{1}{3n^5} \right|$  kann nie größer werden als das positive Argument  $\frac{1}{3} + \frac{2}{n^3} + \frac{1}{3n^5}$ . Also ist:

$$\left| \frac{1}{3} - \frac{2}{n^3} - \frac{1}{3n^5} \right| \leq \frac{1}{3} + \frac{2}{n^3} + \frac{1}{3n^5}$$

- $\frac{1}{3} + \frac{2}{n^3} + \frac{1}{3n^5}$  kann für  $n_0 = 1, n \geq n_0$  nie größer werden als für  $n = n_0$ . Also gilt:

$$\left| \frac{1}{3} - \frac{2}{n^3} - \frac{1}{3n^5} \right| \leq \frac{1}{3} + \frac{2}{n^3} + \frac{1}{3n^5} \leq \frac{8}{3}$$

31.25.2. *Ungleichung mit  $K$  und  $n_0$ .* Quelle: [10]. Berechnen Sie ein  $K$  und ein  $n_0$  mit  $|3n^2 - n - 5| \leq Kn^2$  für alle  $n \geq n_0$ . Da es hier um eine Folge  $(a_n)$  geht, ist sicher  $n_0 \geq 1$ . Wir wählen in diesem Bereich  $n_0$  frei zu  $n_0 = 1$ .

$$\begin{aligned} |3n^2 - n - 5| &\stackrel{!}{\leq} Kn^2 \\ \Rightarrow n^2 \left| 3 - \frac{1}{n} - \frac{5}{n^2} \right| &\stackrel{!}{\leq} Kn^2 \\ \Leftrightarrow \left| 3 - \frac{1}{n} - \frac{5}{n^2} \right| &\leq 3 + \frac{1}{n} + \frac{5}{n^2} \leq 9 \stackrel{!}{\leq} K \end{aligned}$$

Die Gleichung  $9 \stackrel{!}{\leq} K$  ist für  $K = 9$  erfüllt. Die Lösung ist also  $K = 9$  für  $n \geq n_0 = 1$ .

## 31.26. Differentialrechnung.

31.26.1. *Nullstellen, Polstellen und Lücken.* Quelle: [9]. Berechnen Sie für die folgende Funktion die Nullstellen, Polstellen und Lücken der Art  $\frac{0}{0}$ , wobei letztere auf Behebbarkeit zu prüfen sind:  $y = \frac{f(x)}{g(x)} = \frac{x(x+1)(x+3)}{(x+3)^2(x-1)}$ .

**Nullstellen:** Bedingung:  $f(x_N) = 0 \wedge g(x_N) \neq 0$

$$x_{N_1} = 0$$

$$x_{N_2} = -1$$

**Polstellen:** Bedingung:  $f(x_P) \neq 0 \wedge g(x_P) = 0$

$$x_{P_1} = 1$$

**Definitionslücken:** Bedingung  $f(x_D) = 0 \wedge g(x_D) = 0$

$$x_{D_1} = -3$$

Da für alle  $x \neq -3$  gilt

$$y = \frac{x(x+1)(x+3)}{(x+3)^2(x-1)} = \frac{x(x+1)}{(x+3)(x-1)} = y^*$$

ist  $x_{D_1} = -3$  eine mit einer Polstelle behebbare Definitionslücke.

31.26.2. *Nullstellen, Polstellen und Lücken.* Quelle: [10]. Berechnen Sie für die folgende Funktion die Nullstellen, Polstellen und Lücken der Art  $\frac{0}{0}$ , wobei letztere auf Behebbarkeit zu prüfen sind:  $y = \frac{f(x)}{g(x)} = \frac{x(x-1)(x+2)}{(x-1)^2(x-2)}$ .

**Nullstellen:** Bedingung:  $f(x_N) = 0 \wedge g(x_N) \neq 0$

$$x_{N_1} = 0$$

$$x_{N_2} = -2$$

**Polstellen:** Bedingung:  $f(x_P) \neq 0 \wedge g(x_P) = 0$

$$x_{P_1} = 2$$

**Definitionslücken:** Bedingung  $f(x_D) = 0 \wedge g(x_D) = 0$

$$x_{D_1} = 1$$

Da für alle  $x \neq 1$  gilt

$$y = \frac{x(x-1)(x+2)}{(x-1)^2(x-2)} = \frac{x(x+2)}{(x-1)(x-2)} = y^*$$

ist  $x_{D_1} = 1$  eine mit einer Polstelle behebbare Definitionslücke.

31.26.3. *Grenzwert  $\lim_{n \rightarrow \infty}$  einer gebrochenrationalen Funktion.* Quelle: [9]. Berechnen Sie:

$$\begin{aligned} \lim_{n \rightarrow \infty} \frac{(n - n^2)(3n - 4)}{(n + 2)(n - 1)^2} &= \lim_{n \rightarrow \infty} \frac{-3n^3 + 7n^2 - 4n}{n^3 - 3n + 2} \\ &= \lim_{n \rightarrow \infty} \frac{-3 + \frac{7}{n} - \frac{4}{n^2}}{1 - \frac{3}{n^2} + \frac{2}{n^3}} \\ &= -3 \end{aligned}$$

31.26.4. *Grenzwert  $\lim_{n \rightarrow \infty}$  einer gebrochenrationalen Funktion.* Quelle: [10]. Berechnen Sie. Hier bietet es sich an, die Ausdrücke nicht auszurechnen, sondern aus jedem der drei Faktoren in Zähler und Nenner je ein  $n$  auszuklammern. Nach Kürzen sind dann bereits wie gewünscht alle  $n$  im Nenner von Brüchen vorhanden und man kann den Grenzwert bilden, ohne die Ausdrücke jemals vollständig ausgerechnet zu haben.

$$\begin{aligned} \lim_{n \rightarrow \infty} \frac{(1 - \frac{n}{2})^2(1 - 4n)}{(3n - 2)n(2 - \frac{n}{3})} &= \lim_{n \rightarrow \infty} \frac{n(\frac{1}{n} - \frac{1}{2})n(\frac{1}{n} - \frac{1}{2})n(\frac{1}{n} - 4)}{n(3 - \frac{2}{n})n(1)n(\frac{2}{n} - \frac{1}{3})} \\ &= \frac{(-\frac{1}{2})(-\frac{1}{2})(-4)}{3(-\frac{1}{3})} = \frac{-1}{-1} \\ &= 1 \end{aligned}$$

31.26.5. *Ableitung über Differenzenquotient.* Quelle: [9]. Berechnen Sie die erste Ableitung der Funktion  $y = 3x^2 - 2$ , indem Sie den Differenzenquotienten  $\frac{\Delta y}{\Delta x} = \frac{f(x+\Delta x) - f(x)}{\Delta x}$  bilden und nach geeigneter Umformung den Grenzübergang für  $\Delta x \rightarrow 0$  durchführen.

$$\begin{aligned} f'(x) &= \lim_{\Delta x \rightarrow 0} \frac{\Delta y}{\Delta x} \\ &= \lim_{\Delta x \rightarrow 0} \frac{f(x + \Delta x) - f(x)}{\Delta x} \\ &= \lim_{\Delta x \rightarrow 0} \frac{3x^2 + 6x\Delta x + 3\Delta x^2 - 2 - 3x^2 + 2}{\Delta x} \\ &= \lim_{\Delta x \rightarrow 0} \frac{6x\Delta x + 3\Delta x^2}{\Delta x} \\ &= \lim_{\Delta x \rightarrow 0} (6x + 3\Delta x) \\ &= 6x \end{aligned}$$

31.26.6. *Ableitung einer Logarithmusfunktion.* Quelle: [9]. Differenzieren Sie:  $y = (x - 1)^{\ln x}$ .

$$\begin{aligned} y &= (x - 1)^{\ln x} \\ &= e^{\ln((x-1)^{\ln x})} \\ &= e^{\ln x \cdot \ln(x-1)} \end{aligned}$$

$$\begin{aligned} y' &= e^{\ln x \cdot \ln(x-1)} \cdot (\ln x \cdot \ln(x-1))' \\ &= e^{\ln x \cdot \ln(x-1)} \cdot \left( \frac{1}{x} \ln(x-1) + \ln x \cdot \frac{1}{x-1} \right) \\ &= \frac{1}{x} \ln(x-1) \cdot (x-1)^{\ln x} + \ln x \cdot (x-1)^{\ln x-1} \end{aligned}$$

31.26.7. *Grenzwert mit L'Hospital.* Quelle: [9]. Berechnen Sie mittels L'Hospital:

$$\begin{aligned} \lim_{x \rightarrow \pi} \frac{1 + \cos x}{\sin x} &= \lim_{x \rightarrow \pi} \frac{-\sin x}{\cos x} \\ &= \lim_{x \rightarrow \pi} (-\tan x) \\ &= 0 \end{aligned}$$

31.26.8. *Grenzwert mit L'Hospital.* Quelle: [10]. Berechnen Sie mittels L'Hospital:

$$\begin{aligned} \lim_{x \rightarrow 1} \frac{x^3 - 2x + 1}{2x^3 - \frac{1}{2}x^2 - \frac{3}{2}} &= \lim_{x \rightarrow 1} \frac{3x^2 - 2}{6x^2 - x} \\ &= \frac{1}{5} \end{aligned}$$

31.26.9. *Wendepunkt und Krümmungsverhalten einer Funktion.* Quelle: [9]. Gegeben sei  $y = \frac{1}{4}x^4 - 2x^3 + \frac{9}{2}x^2 - 1$ .  
1. Berechnen Sie die Wendepunkte der Funktion.

$$y' = x^3 - 6x^2 + 9x$$

$$y'' = 3x^2 - 12x + 9$$

$$y''' = 6x - 12$$

Notwendige Bedingung:

$$\begin{aligned} y'' &= 0 \\ \Rightarrow 3x^2 - 12x + 9 &= 0 \\ \Leftrightarrow x^2 - 4x + 3 &= 0 \\ \Leftrightarrow (x-2)^2 &= 1 \\ x_{W_1} = 3 \quad \vee \quad x_{W_2} = 1 \end{aligned}$$

Hinreichende Bedingung:

$$y''' \neq 0$$

$$y'''(W_1) = 6 \Rightarrow W_1 \left( 3, \frac{23}{4} \right)$$

$$y'''(W_2) = -6 \Rightarrow W_2 \left( 1, \frac{7}{4} \right)$$

Berechnen Sie das Krümmungsverhalten für  $x = -1$ .

$$y''(-1) = \frac{23}{4} > 0$$

Die Kurve  $y$  ist in  $x = -1$  linksgekrümmt, da  $y''(-1) > 0$ .

31.26.10. *Wendepunkte und Krümmungsverhalten einer Funktion.* Quelle: [10]. Berechnen Sie die Wendepunkte (also  $x$ - und  $y$ -Werte) und untersuchen Sie das Krümmungsverhalten der Funktion

$$f(x) = \frac{1}{6}x^4 - x^3 + 2x^2 + 1$$

$$f'(x) = \frac{2}{3}x^3 - 3x^2 + 4x$$

$$f''(x) = 2x^2 - 6x + 4$$

$$f'''(x) = 4x - 6$$

Notwendige Bedingung:

$$\begin{aligned} f''(x) &= 0 \\ \Rightarrow 2x^2 - 6x + 4 &= 0 \\ \Leftrightarrow x^2 - 3x + 2 + \frac{1}{4} &= \frac{1}{4} \\ \Leftrightarrow \left(x - \frac{3}{2}\right)^2 &= \frac{1}{4} \\ \Leftrightarrow x_1 = 1 \quad \vee \quad x_2 = 2 \end{aligned}$$

Hinreichende Bedingung:

$$\begin{aligned} f'''(x) &\stackrel{!}{\neq} 0 \\ f'''(x_1) &= -2 \\ f'''(x_2) &= 2 \end{aligned}$$

Also sind die Wendepunkte der Funktion:

$$\begin{aligned} W_1(x_1 \mid f(x_1)) &= W_1\left(1 \mid \frac{13}{6}\right) \\ W_2(x_2 \mid f(x_2)) &= W_2\left(2 \mid \frac{11}{3}\right) \end{aligned}$$

Krümmungsverhalten der Funktion:  $0 < x_1 \wedge f''(0) = 1 > 0 \Rightarrow$

- $f(x)$  ist in  $(-\infty; x_1) = (-\infty; 1)$  linksgekrümmt
- $f(x)$  ist in  $(x_1; x_2) = (1; 2)$  rechtsgekrümmt
- $f(x)$  ist in  $(x_2; +\infty) = (2; +\infty)$  linksgekrümmt

An Wendepunkten ändert sich die Krümmung. Deshalb haben wir nur einen Wert  $0 < x_1$  einsetzen müssen und nicht einen aus jedem der drei Teilintervalle.

31.26.11. *Schräge Asymptote.* Quelle: [10]. Berechnen Sie für die folgende Funktion die schräge Asymptote (sofern vorhanden):

$$y = \frac{2x^3 - 2x^2 + 2}{x^2 - 4}$$

Der Grad des Zählers  $n = 3$  ist gegenüber dem Grad des Nenners  $m = 2$ :  $n = m + 1$ . Also existiert eine Gerade als schräge Asymptote, deren Gleichung sich aus dem ganzzahligen Anteil der Polynomdivision ergibt:

$$(2x^3 - 2x^2 + 2) : (x^2 - 4) = 2x - 2 + \frac{8x - 6}{x^2 - 4}$$

$$\Rightarrow y_g = 2x - 2$$

31.26.12. *Differenzieren einer Funktion.* Quelle: [10]. Differenzieren Sie:

$$\begin{aligned} y &= (\ln(\sin x))^2 = \ln(\sin x) \cdot \ln(\sin x) \\ y' &= \ln(\sin x) (\ln(\sin x))' + \ln(\sin x) (\ln(\sin x))' \\ &= 2 \ln(\sin x) (\ln(\sin x))' \\ &= 2 \ln(\sin x) \frac{1}{\sin x} \cos x \\ &= 2 \cdot \ln(\sin x) \cdot \cot(x) \end{aligned}$$

31.26.13. *Differenzieren einer Funktion.* Quelle: [10]. Differenzieren Sie. Hier betrachtet man die Funktion als Gleichung und führt mit ihr Äquivalenzumformungen aus, bis sie beide Seiten einfach differenzierbar sind. Beachte:  $\ln y$  muss nach der Kettenregel abgeleitet werden, denn  $y$  ist selbst ein Ausdruck mit  $x$ .

$$\begin{aligned} y &= (\sin x)^{x-1} \\ \Leftrightarrow \ln y &= \ln \left( (\sin x)^{x-1} \right) = (x-1) \cdot \ln(\sin x) \\ \Rightarrow (\ln y)' &= ((x-1) \cdot \ln(\sin x))' \\ \Leftrightarrow \frac{1}{y} \cdot y' &= \ln(\sin x) + (x-1) \cot x \quad | \cdot y = (\sin x)^{x-1} \\ \Leftrightarrow y' &= (\sin x)^{x-1} (\ln(\sin x) + (x-1) \cot x) \end{aligned}$$

### 31.27. Integralrechnung.

31.27.1. *Integral mit Produktintegration.* Quelle: [9]. Berechnen Sie:

$$\begin{aligned} \int x^2 e^{2x} dx &= \left[ \frac{1}{2} e^{2x} x^2 + C \right] - \int \frac{1}{2} e^{2x} \cdot 2x dx \\ &= \left[ \frac{1}{2} e^{2x} x^2 + C \right] - \left( \left[ \frac{1}{2} e^{2x} x + C \right] - \int \frac{1}{2} e^{2x} dx \right) \\ &= \left[ \frac{1}{2} e^{2x} x^2 + C - \frac{1}{2} e^{2x} x \right] + \frac{1}{2} \int e^{2x} dx \\ &= \left[ \frac{1}{2} e^{2x} x^2 - \frac{1}{2} e^{2x} x + \frac{1}{4} e^{2x} + C \right] \\ &= \frac{1}{2} e^{2x} \left( x^2 - x + \frac{1}{2} \right) + C \end{aligned}$$

### LITERATUR

- [1] Skript zur Vorlesung Mathematik 1 von Prof. Eichner, Wintersemester 1996/1997. Studentische Mitschrift - zum Download als zip-Datei (19.699 Byte), entpackt als doc-Datei (113.664 Byte). Erhältlich unter Skriptsammlung der Fachschaft MNI der FH Gießen-Friedberg <http://www.fh-giessen.de/FACHSCHAFT/Informatik/cgi-bin/navi01.cgi?skripte>. In der Veranstaltung im WS 2002/2003 wurde die Aussagenlogik etwas intensiver behandelt als in dieser Mitschrift dargestellt. Evtl. gibt es eine weitere Mitschrift auf Papier in der Fachschaft auszuleihen.
- [2] Marcus Martin: »Brückenkurs in Mathematik (BWL)«. Quelle früher <http://www.uni-giessen.de/~gc1103/vorkfh.htm>, jetzt unbekannt. Deshalb zur Verfügung gestellt unter <http://matthias.ansorgs.de/InformatikDiplom/Modul.VkMa.Fremdt/BrueckenkursMathe.zip>. Der Inhalt der Datei `a0.pdf` wurde vollständig in das vorliegende Dokument integriert.
- [3] Dokumentation zum Aufbau von float-Datentypen siehe `/usr/include/ieee754.h` auf \*X-Systemen.
- [4] Papula: »Mathematik für Fachhochschulen«. Empfehlenswert zur Begleitung der Vorlesung Mathematik 1, um den Stoff in der Litertur nachzuvollziehen.
- [5] Papula: »Mathematik für Ingenieure und Naturwissenschaftler«; 3 Bände; Vieweg-Verlag. Diese Mathematikbücher werden von Studenten bevorzugt, jedoch nicht von Prof. Eichner.
- [6] Hartmann: »Mathematik für Informatiker«; Vieweg-Verlag 2002. Dieses Buch wird von Prof. Eichner bevorzugt.
- [7] Fetzer und Fränkel: »Mathematik«; 2 Bände; VDI-Verlag. Dieses Buch wird von Prof. Eichner bevorzugt.
- [8] Meinel und Mundhenk: »Mathematische Grundlagen der Informatik«; 2 Auflage 2002.
- [9] Prof. Dr. L. Eichner: »Fach: Mathematik1; Klausur am 16.7.2001«; Fachbereich MNI; FH Gießen-Friedberg.
- [10] Prof. Dr. L. Eichner: »Fach: Mathematik1; Klausur am 28.1.2000«; Fachbereich MNI; FH Gießen-Friedberg.
- [11] Joachim Ansorg: »Abiturvorbereitung Mathematik«; »Abiturvorbereitung 2002 für den Mathekurs von Herrn Heinrichs (1999-2001) am Gymnasium Marienstatt (Rheinland-Pfalz)«. Im Internet unter <http://joachim.ansorgs.de>. Genutzt mit Erlaubnis. Übernommen wurden alle Kapitel außer »Trigonometrie« und allen Kapiteln aus »Teil 3. Stochastik«.
- [12] Thimotheus Pokorra: »Mathematik I; WS 98/99 bei Löffler«. Quelle Homepage von Thimotheus Pokorra <http://homepages.fh-giessen.de/~hg9541/mitschr.html>, direkter Link <http://homepages.fh-giessen.de/~hg9541/mathe.zip>. Später wohl auf <http://www.pokorra.de>.
- [13] Prof. Dr. Löffler: Aufgabenblätter zur Vorlesung Mathematik I im Wintersemester 2001/2002. Quelle: Homepage von Rolf H. Viehmann <http://www.rolfhub.de/de/study.ws0102.phtml>.
- [14] Tim Pommerening: »Mathe 1 Formelsammlung: Vektoren - 1. Semester; (WS 00/01 Prof. Dr. Metz)«. Referenziert auf Homepage von Tim Pommerening: Bereich Dokumente <http://homepages.fh-giessen.de/~hg12061/dokumente.html>, direkter Link <http://homepages.fh-giessen.de/~hg12061/docs/vektoren.rar>. Ein sehr nützliches Dokument!
- [15] Tobias Webelsiep: »Mathematik I; Bei Herrn Metz«; FH Gießen-Friedberg; WS 2000. Referenziert auf [http://www.webelsiep.de/studium\\_skripte.php](http://www.webelsiep.de/studium_skripte.php) bzw. über die Webelsiep-Homepage <http://www.webelsiep.de/default.php?main=studium>. Direkter Link <http://www.webelsiep.de/downloads/skripte/mathematik1.pdf>. Ein wirklich brauchbares Dokument, nur der Ausdruck mit Acrobat Reader unter Linux kann problematisch werden.