



universität
wien

DIPLOMARBEIT / DIPLOMA THESIS

Titel der Diplomarbeit / Title of the Diploma Thesis

Konstruktion der Zahlenbereiche

Von den Axiomen der Mengenlehre über unvernünftige und eingebildete Zahlen
zu den Oktonionen

verfasst von / submitted by

Lukas Weissenböck

angestrebter akademischer Grad / in partial fulfilment of the requirements for the degree of

Magister der Naturwissenschaften (Mag. rer. nat.)

Wien, 2018 / Vienna, 2018

Studienkennzahl lt. Studienblatt / degree programme code as it appears on the student record sheet: A 190 406 884

Studienrichtung lt. Studienblatt / degree programme as it appears on the student record sheet: Lehramtsstudium UF Mathematik UF Informatik und Informatikmanagement

Betreut von / Supervisor: ao. Univ.-Prof. Mag. Dr. Peter Raith

No one knows what the future holds. That's why its potential is infinite.

— Okabe Rintarou, *Steins;Gate*

Inhaltsverzeichnis

Einleitung	1
1 Die natürlichen Zahlen	5
1.1 Ursprung der natürlichen Zahlen	5
1.2 Axiomatische Mengenlehre	9
1.2.1 Die Axiome von Zermelo und Fraenkel	11
1.3 Konstruktion der natürlichen Zahlen	14
1.4 Rechenoperationen auf $\mathbb{N}$	18
1.5 Die natürliche Ordnung	29
1.5.1 Die natürliche Ordnung als Wohlordnung	33
2 Die ganzen Zahlen	35
2.1 Geschichte der negativen Zahlen	36
2.2 Konstruktion der ganzen Zahlen	37
2.3 Rechenoperationen auf $\mathbb{Z}$	41
2.4 Einbettung von $\mathbb{N}$ in $\mathbb{Z}$	53
2.5 Der Hauptsatz der elementaren Zahlentheorie	55
3 Die rationalen Zahlen	61
3.1 Konstruktion der rationalen Zahlen	61
3.2 Rechenoperationen auf $\mathbb{Q}$	65
3.3 Einbettung von $\mathbb{Z}$ in $\mathbb{Q}$	75
3.4 $\mathbb{Q}$ als kleinster Körper über $\mathbb{Z}$	79
4 Die reellen Zahlen	83
4.1 Die Entdeckung der Inkommensurabilität	84
4.2 LÖcher im Zahlenstrahl	87
4.3 Konstruktion der reellen Zahlen	91
4.4 Rechenoperationen auf $\mathbb{R}$	95
4.5 Vollständigkeit von $\mathbb{R}$	104
4.6 Einbettung von $\mathbb{Q}$ in $\mathbb{R}$	109
4.7 Unendlichkeit der Zahlenmengen	111

5	Die komplexen Zahlen	117
5.1	Lösungen von Gleichungen höheren Grades	118
5.2	Konstruktion der komplexen Zahlen	120
5.3	Rechenoperationen auf $\mathbb{C}$	121
5.4	Einbettung von $\mathbb{R}$ in $\mathbb{C}$	125
5.5	Polardarstellung komplexer Zahlen	127
5.6	Der Fundamentalsatz der Algebra	130
6	Die Quaternionen	137
6.1	Konstruktion der Quaternionen	138
6.2	Rechenoperationen auf $\mathbb{H}$	140
6.3	Einbettung von $\mathbb{C}$ in $\mathbb{H}$	144
6.4	Divisionsalgebren im $\mathbb{R}^{2 \cdot n+1}$, Satz von Frobenius	147
7	Die Oktonionen	151
7.1	Konstruktion der Oktonionen	151
7.2	Rechenoperationen auf $\mathbb{O}$	152
7.3	Einbettung von $\mathbb{H}$ in $\mathbb{O}$	159
	Bemerkungen zur Didaktik der Zahlenbereichserweiterungen	163
	Ausblick	167
	Zusammenfassung	171
	Abstract	173
	Abbildungsverzeichnis	175
	Literatur	177

Einleitung

Die Zahlengerade. Das hier ist Null, Eins, Zwei, Drei – die natürlichen Zahlen. Es gibt ganze Zahlen, da sind auch die negativen dabei. Es gibt Brüche, die rationalen Zahlen, es gibt Wurzeln, die irrationalen Zahlen. Immer ergeben sich neue Zahlen aus Rechenaufgaben; immer muss der Bereich erweitert werden.

C. F. Gauß, Die Vermessung der Welt (Film, 2012)

Zahlen zählen zu den größten Entdeckungen der Menschheit. Es ist unklar, was unsere steinzeitlichen Vorfahren dazu bewegt hat, mit dem Zählen zu beginnen – dass sich aus ihren gekerbten Knochen eine der ertragreichsten Wissenschaften entwickeln würde, wäre ihnen allerdings vermutlich niemals eingefallen.

Thema dieser Diplomarbeit sind Zahlen, oder besser gesagt: die verschiedenen Zahlenbereiche, von denen einige im vorigen Zitat genannt wurden. Die natürlichen Zahlen, mit denen alles seinen Anfang nahm, reichten einige tausend Jahre später für die ersten Hochkulturen nicht mehr zur Lösung praktischer Probleme aus, und wurden auf die rationalen Zahlen erweitert. Im antiken Griechenland erkannte man, dass auch die rationalen Zahlen – im wahrsten Sinne des Wortes – nicht vollständig sind, und erweiterte sie um die irrationalen Zahlen. Ein paar Jahrhunderte später wurden in China zur Lösung neuer Probleme erstmals negative Zahlen verwendet, während man im 17. Jahrhundert erneut an die Grenzen der bisherigen Zahlenbereiche stieß, und so zu den komplexen Zahlen gelangte.

Zu jedem Zahlenbereich in dieser Arbeit wollen wir unter anderem folgende Fragen beantworten:

- Wie lässt sich der Zahlenbereich aus dem vorherigen konstruieren, sodass wir möglichst viele Eigenschaften mitnehmen können (Operatoren, Notation, ...)?
- Welche Gründe gab es historisch für die Erweiterung?
- Welche Eigenschaften besitzt der Zahlenbereich?
- Wodurch zeichnet sich der Zahlenbereich von anderen ab?

In Kapitel 1 konstruieren wir die natürlichen Zahlen aus den Axiomen der Mengenlehre nach Zermelo und Fraenkel. Wir starten mit einem Abriss über den Ursprung der natürlichen Zahlen, gehen weiter zur Mengenlehre, bis wir die natürlichen Zahlen und bekannte Operationen auf ihnen definieren. Wir schließen das Kapitel mit einem Beweis der Wohlordbarkeit der natürlichen Zahlen.

In Kapitel 2 führen wir die erste Zahlenbereichserweiterung auf die ganzen Zahlen durch. Wir beginnen wieder mit einem historischen Überblick, und widmen uns im Anschluss der Definition bekannter Operationen. In Abschnitt 2.4 werden wir der Frage nachgehen, was eine Zahlenbereichserweiterung formal gesehen ist. Wir beenden das Kapitel mit einem Beweis des Hauptsatzes der elementaren Zahlentheorie.

Kapitel 3 widmet sich den rationalen Zahlen, wobei wir analog zu den vorherigen Kapiteln vorgehen. Das Hauptresultat dieses Kapitels wird sein, dass die rationalen Zahlen den kleinsten Körper über den ganzen Zahlen bilden.

Das arbeitsintensivste Kapitel dieser Arbeit stellt Kapitel 4 dar, in welchem wir die reellen Zahlen konstruieren. Zu Beginn zeigen wir, wie man im antiken Griechenland erkannt hat, dass es Zahlen gibt, die sich nicht durch Brüche ausdrücken lassen. Wir formalisieren das Problem der rationalen Zahlen mit Hilfe von Intervallschachtelungen und Cauchyfolgen, und gelangen so zum Begriff der Vollständigkeit und zur Menge der reellen Zahlen. Zum Schluss des Kapitels beschäftigen wir uns mit der Größe der Zahlenbereiche, und erhalten so einige bekannte, unserer Intuition widersprechende Resultate über unendlich große Mengen.

In Kapitel 5 erweitern wir, auf der Suche nach Nullstellen von Polynomen, die reellen Zahlen auf die komplexen Zahlen. Wieder werden wir Antworten auf die zuvor gestellten

Fragen geben, wobei das Hauptresultat dieses Kapitels ein Beweis des Fundamentalsatzes der Algebra ist.

Auf eine exotischere Zahlenmenge führt Kapitel 6, in welchem wir die komplexen Zahlen auf die Quaternionen erweitern. Wir werden in diesem Kapitel erstmals schmerzlich feststellen, dass Zahlenbereichserweiterungen stets einen Preis haben – so ist etwa die Multiplikation in den Quaternionen nicht mehr kommutativ. In Abschnitt 6.4 werden wir mit Hilfe des Satzes von Frobenius und eines weiteren Resultats zeigen, dass Zahlenbereichserweiterungen bestimmte Grenzen gesetzt sind.

Noch bizarrer wird Kapitel 7, in welchem wir die Oktonionen als achtdimensionale Zahlen erhalten – allerdings zu einem hohen Preis: Die Multiplikation ist hier weder kommutativ noch assoziativ, was den Umgang mit Oktonionen sehr mühsam macht.

Die Arbeit schließt mit einigen Bemerkungen zur Didaktik der Zahlenbereichserweiterungen. Wir fassen zudem einige weitere Resultate über Zahlenbereichserweiterungen zusammen und skizzieren eine Beweisidee für den Satz von Hurwitz über endlichdimensionale reelle Divisionsalgebren.

Eine ausführliche Übersicht über verwendete Quellen befindet sich am Ende dieser Arbeit. Wo immer ich bei Resultaten, Formulierungen von Sätzen oder Beweisideen Fremdquellen verwendet habe, wird das direkt im Text, oder bei Sätzen in der Klammer nach der Bezeichnung, durch Angabe der Quelle deutlich. Separat möchte ich an dieser Stelle hervorheben, dass einige Ideen zur Definition der natürlichen Zahlen aus dem Buch *Einführung in das mathematische Arbeiten* von Univ.-Prof. Dr. Hermann Schichl und Univ.-Prof. Dr. Roland Steinbauer stammen, während ich Grundideen für die Definition der reellen Zahlen aus der Vorlesung *Einführung in die Analysis*, gehalten von Dr. Bernhard Krön im Sommersemester 2013 an der Universität Wien, habe. Einige Konzepte aus der Vorlesung *Zahlenbereichserweiterungen*, gehalten von Univ.-Prof. Dr. Günther Karigl im Wintersemester 2017/18 an der Technischen Universität Wien, sind ebenfalls in dieser Arbeit verwebt.

An vielen Stellen werden Begriffe und Resultate der Algebra und Analysis als bekannt vorausgesetzt, und nur dann bewiesen oder hervorgehoben, wenn sie für einen Zahlenbereich von Interesse sind. Das betrifft unter anderem Rechenregeln in Ringen, die Potenzschreibweise, aber auch Definitionen und Eigenschaften von Grenzwerten.

Danksagungen

Vielen lieben Dank Dir, Peter – für die zahlreichen interessanten Lehrveranstaltungen im Verlauf meines Studiums, und dass ich zu diesem Thema meine Diplomarbeit bei Dir schreiben konnte, die Du ausgezeichnet betreut hast.

Danke auch meinen Eltern für die Unterstützung während des Studiums und der vergangenen 25 Jahre.

1 Die natürlichen Zahlen

Die natürlichen Zahlen, also die bekannten Zählzahlen *Null, Eins, Zwei, Drei, ...*, sind in mehrererlei Hinsicht *der* grundlegende Zahlenbereich. Zum einen lassen sich von ihnen ausgehend alle wichtigen Zahlenbereiche der Mathematik konstruieren, wie es in dieser Diplomarbeit gezeigt wird. Zum anderen sind die natürlichen Zahlen historisch gesehen der erste Zahlenbereich, der von Menschen verwendet wurde. Auch in der Entwicklung von Kindern spielen das Erkennen von Anzahlen und die natürlichen Zahlen in Form der Zählzahlen eine herausragende Rolle (vgl. dazu [51, S. 344–345; 10, S. 118]).

1.1 Ursprung der natürlichen Zahlen

Es ist unbekannt, wann Menschen zum ersten Mal gezählt haben. Es gibt allerdings Hinweise darauf, dass der Mensch in seiner Entwicklung eine lange Periode durchlief, in der er nicht zählen konnte [31, S. 23–27]. Als Vorläufer des Zählens findet sich nach Dantzig beim Menschen, wie auch bei vielen Tieren, ein Gefühl für Anzahlen (*number sense*) [17, S. 1]. Dieses Gefühl erlaubte es dem Menschen eine Menge auf intuitive Weise zahlenmäßig zu erfassen, wodurch er dazu in der Lage war zu erkennen, ob einer kleinen Menge ein Element hinzugefügt oder weggenommen wurde. Auch einige Tierarten, wie etwa Vögel, verfügen über ein solches Gefühl, und können so beispielsweise feststellen, ob eine kleine Futtermenge größer als eine andere ist [31, S. 21].

Das Gefühl für Anzahlen ist bei Tieren, wie auch bei Menschen, sehr beschränkt. Mengen von bis zu vier Elementen können relativ problemlos voneinander unterschieden werden, ab fünf oder mehr Elementen wird eine intuitive Erfassung jedoch schwierig. Als Beispiel für diese Schwierigkeit nennt Ifrah verschiedene Stämme von Eingeborenen in Afrika

und Ozeanien, in deren Sprachen für Anzahlen ab vier Elementen Ausdrücke wie „viel“ oder „unzählig“ verwendet werden [31, S. 25]. Dantzig spricht die Doppeldeutigkeit von „thrice“ im Englischen an, das sowohl für „drei“, als auch für „viel“ steht [17, S. 5]. Auch die grammatikalischen Zahlformen (Singular, Dual, Plural) weisen auf diese Grenze hin.

Die intuitive Zahlenwahrnehmung hat sich bis heute kaum verbessert. So gruppieren wir beispielsweise Strichlisten in Fünfergruppen, wobei der fünfte Strich zur besseren Übersicht üblicherweise nicht neben die anderen vier, sondern quer darüber gesetzt wird.

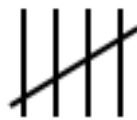


Abbildung 1.1: Der fünfte Strich wird quer über die anderen vier gesetzt.

Da das intuitive Gefühl für Anzahlen auch in der Entwicklungspsychologie eine wichtige Rolle spielt, geht man davon aus, dass sich das Konzept des Zählens aus diesem entwickelt hat [51, S. 344; 17, S. 5].

Die frühesten Funde, die menschliches Zählen belegen, sind zwischen 20 000 und 30 000 Jahre alt, und liegen in Form von Einkerbungen in Knochen oder an Felswänden neben Zeichnungen vor [31, S. 14]. In Bilzingsleben wurden von *Homo erectus* absichtsvoll gekerbte Knochen gefunden, die etwa 370 000 Jahre alt sind – deren Deutung als Hilfsmittel fürs Zählen ist aber nicht belegt [56, S. 12].

Der Mensch dürfte mit dem Zählen begonnen haben um feststellen zu können, ob Anzahlen gleich oder unterschiedlich sind [56, S. 7]. Für solche Vergleiche war noch keine abstrakte Zähltechnik notwendig, wie wir sie heute verwenden, nämlich unter

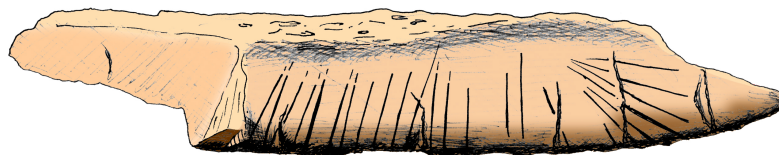


Abbildung 1.2: Gekerbter Knochen, etwa 370 000 Jahre alt [14].

Verwendung abstrakter Zahlwörter. Ifrah nennt als Beispiel für ein konkretes Zählverfahren einen Hirten, der wissen möchte, ob Abends all seine Tiere vollzählig sind. Dieser könnte am Morgen für jedes Tier, das an ihm vorbeigeht, eine Kerbe in einen Knochen schnitzen. Am Abend würde er für jedes Tier, das er sieht, mit seinem Finger eine Kerbe weiterwandern. Ist er mit dem letzten Tier gleichzeitig auch bei der letzten Kerbe angekommen, wäre die Herde vollzählig [31, S. 28].

Die grundsätzliche Idee dieses konkreten Zählverfahrens ist, für das Zählen zunächst eine Hilfsmenge heranzuziehen (Steine, Kerben, ...). Anschließend werden paarweise aus der zu zählenden Menge und der Hilfsmenge Elemente entfernt. Sind am Ende in der Hilfsmenge noch Elemente übrig, war diese ursprünglich größer als die zu zählende Menge. Reichen umgekehrt die Elemente der Hilfsmenge für eine paarweise Zuordnung nicht aus, war die zu zählende Menge größer als die Hilfsmenge. Falls sich jedoch paarweise alle zu zählenden Elemente der gesamten Hilfsmenge zuordnen lassen, dann sind beide Mengen gleich groß.

Dieses konkrete Zählen mit Hilfsmengen war tatsächlich tausende Jahre in Verwendung, wie ein Fund aus Mesopotamien zeigt. Dort wurde 1928/29 ein Gefäß aus Ton entdeckt, auf dem außen eine Inventarliste aufgeschrieben war. Innen befanden sich entsprechend der außen notierten Gesamtanzahl viele kleine Tonkügelchen [31, S. 29]. Oppenheim schildert eine Episode während der Grabungsarbeiten, durch welche die Bedeutung des Funds als Hilfsmenge zum Zählen klar wurde [41, S. 123]:

A servant of the expedition was sent to market to buy chickens. On his return, owing to some accident, the chicken he had bought in town mingled with those kept in the court yard of the expedition house. Since the servant's arithmetic was somewhat insufficient, it seemed for a moment impossible to settle the accounts of the money he was given for his purchases. But he produced a number of pebbles which he had put aside, one for each chicken he had bought [...].

In der Sprache der Mathematik kann man dieses konkrete Zählverfahren so formulieren: Gegeben seien zwei Mengen A und B . Falls es möglich ist, eine Zuordnung zwischen diesen Mengen zu finden, die jedem Element aus A genau ein Element aus B zuordnet, und jedem Element aus B dann ein Element aus A zugeordnet ist, dann haben die

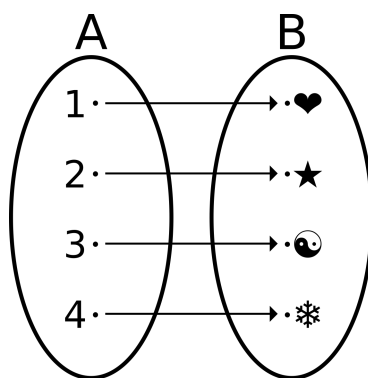


Abbildung 1.3: Beispiel für eine bijektive Zuordnung. Da es eine solche gibt, haben beide Mengen A und B gleich viele Elemente [13].

Mengen gleich viele Elemente. Anders formuliert: Die Mengen A und B sind *gleichmächtig*. Eine Zuordnung, welche diese Eigenschaften erfüllt, wird *bijektiv* genannt (siehe Abbildung 1.3).

Dieses skizzierte Verfahren hat Georg Cantor im 19. Jahrhundert benutzt, um die Gleichmächtigkeit von Mengen zu definieren. Dadurch gelangte er zu seinen erstaunlichen Erkenntnissen über unendlich große Mengen (siehe Abschnitt 4.7). Jene tausende Jahre alte Idee lebt damit in der Basis der modernen Mathematik weiter.

Aus der konkreten Form des Zählens kann sich nun der abstrakte Begriff der Anzahl entwickelt haben. Mit der Zeit wird man bemerkt haben, dass man mit fünf Kerben genauso zählen kann wie mit fünf Steinen oder den Fingern an einer Hand. Es wurde erkannt, dass es eine Eigenschaft gibt, welche diesen verschiedenen Objekten gemeinsam ist: die Anzahl. Diese Erkenntnis war ein wichtiger Schritt hin zum abstrakten Zahlbegriff. [31, S. 35–36].

Dank dieser Erkenntnis reicht es aus, beim konkreten Zählen – anstatt eine Hilfsmenge speziell für diese Situation anzufertigen – bereits im Voraus für jede Anzahl eine bestimmte, naheliegende Hilfsmenge zu definieren, etwa die Hand für „fünf“. Dantzig vermutet, dass diese Einschränkung der Hilfsmengen zur Bildung von Zahlwörtern, und damit letztendlich zum abstrakten Zahlbegriff geführt hat [17, S. 7–8]. Ist es nämlich möglich, mit fünf Fingern, also einer Hand, stets die Anzahl „fünf“ zu erfassen, bietet es sich an, für diese Anzahl eben einen Begriff zu verwenden, der mit „eine Hand“ in Beziehung steht. Unterstützt wird diese Vermutung durch Studien an verschiedenen

Naturvölkern [vgl. 31, S. 40]. So verwenden beispielsweise die Lengua-Indianer in Paraguay für die Zahl *Fünf* den Ausdruck „eine Hand“; für *Zehn* sagen sie „beide Hände voll“, während für *Zwanzig* der Ausdruck „Füße voll“ benutzt wird [31, S. 38].

Die Darstellung von Zahlen ist somit wesentlich älter als der abstrakte Zahlbegriff. Durch die Entwicklung der Schrift änderte sich diese. An die Stelle konkreter Kerben traten abstraktere Darstellungen, wie sie etwa die Ägypter und Sumerer vor rund 5 000 Jahren erstmals verwendeten [31, S. 545; 17, S. 22]. Neu war hier die Idee der *Bündelung*: Zahlen von Eins bis Neun wurden durch Hintereinanderschreiben des selben Symbols dargestellt, ab der Zehn wurde ein neues Symbol verwendet. Die Ägypter bündelten anschließend weitere Potenzen von Zehn (100, 1000, 10 000, ...), während die Sumerer die Sechzig als Ausgangswert für weitere Bündelungen verwendeten: 60 , $600 = 60 \cdot 10$, $3600 = 60^2$ und $36\,000 = 60^2 \cdot 10$ [31, S. 210–231].

Einen Meilenstein in der Darstellung von Zahlen stellt die Entwicklung des ersten Stellenwertsystems durch die Babylonier vor etwa 4 000 Jahren dar. Hier waren für den Wert einer Zahl erstmals die Positionen der Ziffern von Bedeutung, und nicht deren grafische Darstellung. Als Basis für ihr Zahlensystem wählten die Babylonier die Zahl 60. Dieses Sexagesimalsystem war zwar für Berechnungen äußerst leistungsfähig, aufgrund der hohen Basis allerdings sehr unhandlich zu benutzen (das „Kleine 1x1“ besteht aus $60 \cdot 60 = 3600$ Einträgen) [31, S. 411; 56, S. 128–130].

Weitere Zahldarstellungen der Antike sind etwa durch die griechischen Systeme und das römische System gegeben ([vgl. 31, S. 165; 56, S. 153]). Die heute von uns benutzten Ziffern gehen auf die Inder zurück, und gelangten über den arabischen Raum nach Europa [31, S. 540–544].

1.2 Axiomatische Mengenlehre

Die Mengenlehre wurde von Georg Cantor in der zweiten Hälfte des 19. Jahrhunderts entwickelt [19, S. 10]. Im Zentrum der Mengenlehre steht die Untersuchung von *Mengen*, die Cantor folgendermaßen definiert [7, S. 481]:

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

Beispiele für Mengen nach dieser Definition wären etwa die Menge aller geraden Zahlen, die Menge aller Katzen oder auch die Menge aller Ideen. Das letzte Beispiel ist insofern interessant, weil diese Menge sich selbst als Element enthält – denn schließlich ist sie ebenfalls eine Idee. Umgekehrt muss sich nicht jede Menge selbst enthalten: Die Menge aller Katzen ist keine Katze.

Die Frage, welche Objekte zu Mengen zusammengefasst werden dürfen, wurde von Cantor nicht eindeutig beantwortet. Ihm war allerdings klar, dass der Bildung von Mengen Grenzen gesetzt sind, etwa wenn die Zusammenfassung von Elementen auf Paradoxa führt [19, S. 187–189]. Daher entwickelte er bereits 1898 Axiome für die Bildung von Mengen, die er allerdings nur brieflich David Hilbert mitteilte [19, S. 443].

Ein berühmtes Paradoxon der Mengenlehre wurde von Bertrand Russell im Jahr 1902 entdeckt. Russell betrachtete die Menge aller Mengen, die sich nicht selbst enthalten, und stellte die Frage, ob sich diese Menge selbst enthält. Enthält sie sich selbst, muss sie sich laut Definition nicht selbst enthalten. Falls sie sich allerdings nicht selbst enthält, muss sie sich aufgrund der Definition selbst enthalten – Widerspruch [19, S. 184–185].

Das Paradoxon von Russell zeigte, dass genauer definiert werden musste, welche Zusammenfassungen als Mengen betrachtet werden dürfen. Um das Problem zu lösen, entwickelte Russell zwischen 1902 und 1908 seine Typentheorie, die sich jedoch nicht allgemein durchsetzen konnte. Parallel dazu erarbeitete Ernst Zermelo sieben Basisaxiome für die Mengenlehre, die er 1908 veröffentlichte. Diese wurden 1922 von Abraham Fraenkel ergänzt, und von Thoralf Skolem im Jahr 1929 in der Sprache der Prädikatenlogik formuliert. Heute bilden diese Axiome, abgekürzt *ZFC* genannt, die Basis für die Mengenlehre [19, S. 417].

ZFC löst die bekannten Paradoxa auf. Aus dem Aussonderungsaxiom leitet sich beispielsweise ab, dass Mengen wie jene von Russell nicht gebildet werden dürfen. Kann es noch weitere, unentdeckte Widersprüche geben? Ist es möglich zu beweisen, dass es in ZFC keine Widersprüche gibt? Die Antwort auf diese Frage hat Kurt Gödel im Jahr

1931 gegeben: Nach dem zweiten Gödelschen Unvollständigkeitssatz lässt sich die Widerspruchsfreiheit von ZFC nicht innerhalb von ZFC beweisen, vorausgesetzt dass ZFC widerspruchsfrei ist [24, S. 196–197]. Es besteht damit prinzipiell die Möglichkeit, dass auch in ZFC Paradoxa auftreten können. Bis heute wurden allerdings keine entdeckt.

1.2.1 Die Axiome von Zermelo und Fraenkel

Die ursprünglichen sieben Basisaxiome wurden von Zermelo im Jahr 1908 vorgestellt: *Extensionalitätsaxiom* (im Original: „Axiom der Bestimmtheit“ [57, S. 263]), *Leermengenaxiom* und *Paarmengenaxiom* (im Original zusammengefasst unter dem „Axiom der Elementarmengen“ [57, S. 263]), *Aussonderungsaxiom*, *Potenzmengenaxiom*, *Vereinigungsaxiom*, *Auswahlaxiom* und *Unendlichkeitsaxiom*. Diese Axiome wurden 1922 von Fraenkel um das *Ersetzungsaxiom* ergänzt [23, S. 231]. Den Abschluss bildet das von Zermelo im Jahr 1930 hinzugefügte *Fundierungsaxiom* [19, S. 434–436].

Die folgende Darstellung der Axiome beruht auf [19, S. 422–436, 454–456; 50, S. 183–185].

Axiom 1 (Extensionalitätsaxiom). Zwei Mengen X und Y sind genau dann gleich, wenn sie die gleichen Elemente enthalten.

$$\forall X : \forall Y : X = Y \Leftrightarrow \forall Z : (Z \in X \Leftrightarrow Z \in Y)$$

Axiom 2 (Leermengenaxiom). Es gibt eine Menge X , die keine Elemente hat.

$$\exists X : \forall Y : Y \notin X$$

Diese Menge X wird mit $\emptyset$ oder $\{\}$ angeschrieben.

Axiom 3 (Paarmengenaxiom). Für zwei Mengen X und Y gibt es eine Menge Z , die genau X und Y enthält.

$$\forall X : \forall Y : \exists Z : \forall U : (U \in Z \Leftrightarrow (U = X \vee U = Y))$$

Wir notieren die Menge Z als $\{X, Y\}$ und $\{X, X\}$ als $\{X\}$.

Axiom 4 (Vereinigungsaxiom). Für jede Menge X gibt es eine Menge Y , welche genau die Elemente aller Elemente von X enthält.

$$\forall X : \exists Y : \forall Z : (Z \in Y \Leftrightarrow \exists U : (U \in X \wedge Z \in U))$$

Diese Menge Y wird als $\bigcup X$ geschrieben, und ist aufgrund des Extensionalitätsaxioms eindeutig.

Axiom 5 (Aussonderungsaxiom). Das Aussonderungsaxiom erlaubt es, mit Hilfe eines Prädikates φ all jene Elemente einer Menge X auszuwählen, für die φ wahr ist. Es wird also eine Teilmenge Y aus einer gegebenen Menge konstruiert.

$$\forall X : \forall p_1, \dots, p_n : \exists Y : \forall U : (U \in Y \Leftrightarrow (U \in X \wedge \varphi(U, p_1, \dots, p_n)))$$

Die Menge Y wird als $\{U \in X \mid \varphi(U)\}$ angeschrieben.

Axiom 6 (Unendlichkeitsaxiom). Es gibt eine Menge X , welche die leere Menge enthält. Für jedes $Y \in X$ ist auch $Y \cup \{Y\} \in X$.

$$\exists X : \forall Y : (\emptyset \in X \wedge (Y \in X \Rightarrow Y \cup \{Y\} \in X))$$

Die Menge X hat also $\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\}, \{\emptyset, \{\emptyset, \{\emptyset\}\}\}, \dots$ als Elemente. Das Unendlichkeitsaxiom wird bei der Konstruktion der natürlichen Zahlen eine wichtige Rolle spielen.

Axiom 7 (Potenzmengenaxiom). Das Potenzmengenaxiom garantiert zu jeder vorhandenen Menge X die Existenz ihrer Potenzmenge, also der Menge all ihrer Teilmengen. Bezeichnet wird diese mit $\mathcal{P}(X)$.

$$\forall X : \exists \mathcal{P} : \forall P : (P \in \mathcal{P} \Leftrightarrow \forall Y : (Y \in P \Rightarrow Y \in X))$$

Axiom 8 (Auswahlaxiom). Für jede Menge X von nichtleeren, paarweise disjunkten Mengen gibt es eine Menge Y , die aus jeder Menge in X genau ein Element enthält.

$$\begin{aligned} \forall X : (\emptyset \notin X \wedge ((\forall U, V : U \in X \wedge V \in X \wedge U \neq V) \Rightarrow \\ \forall W : (W \in U \Rightarrow W \notin V)) \Rightarrow \exists Y : \forall X : (X \in X \Rightarrow \exists ! Z : (Z \in Y \wedge Z \in X))) \end{aligned}$$

Axiom 9 (Ersetzungsaxiom). Das Ersetzungsaxiom sagt, salopp ausgedrückt, dass für jede Menge X und eine Zuordnung f das Bild von X unter f wieder eine Menge ist. Angeschrieben wird diese neue Menge als $f(X)$ oder $\{f(x) \mid x \in X\}$.

Axiom 10 (Fundierungsaxiom). Jede nichtleere Menge X enthält ein Element Y , welches zu X disjunkt ist.

$$\forall X : (X \neq \emptyset \Rightarrow \exists Y : (Y \in X \wedge \nexists Z : (Z \in X \wedge Z \in Y)))$$

Über die Axiome lassen sich die bekannten Operationen auf Mengen (Durchschnitt, Vereinigung) definieren [19, S. 423–426].

Russells Paradox wird durch ZFC insofern aufgelöst, als da die Axiome, insbesondere das Aussonderungsaxiom, die Definition der „Menge aller Mengen, die sich nicht selbst enthalten“ nicht zulassen. Denn die Definition dieser Menge führt auf den Ausdruck $\{X \mid X \notin X\}$. Das Aussonderungsaxiom erlaubt allerdings nur die Konstruktion einer Teilmenge aus einer *gegebenen* Menge. In der Definition $\{X \mid X \notin X\}$ steht nicht, aus welcher Menge die Elemente X kommen – und damit ist die Definition nicht zulässig [19, S. 426].

Mit Hilfe des Fundierungsaxioms können wir zeigen, dass es keine Menge gibt, die sich selbst enthält. Es gilt sogar allgemeiner, dass es keine beliebig verschachtelte Mengenkette $A_1 \in A_2 \in \dots \in A_n$ mit $A_n \in A_1$ geben kann (der Fall der Menge, die sich selbst enthält, lässt sich auf $n = 1$ zurückführen).

Proposition 1.2.1. *Es gibt keine Mengen $A_1 \in A_2 \in \dots \in A_n$ mit $A_n \in A_1$.*

Beweis. Angenommen, es gäbe eine solche verschachtelte Mengenstruktur. Dann betrachten wir die Menge $X = \{A_1, A_2, \dots, A_n\}$, deren Existenz aus Axiom 3 und Axiom 4 folgt. Diese Menge ist durch ihre Konstruktion nichtleer, weswegen es nach Axiom 10 ein $Y \in X$ geben muss, sodass Y zu X disjunkt ist.

Dieses Y kann nun nicht A_1 sein, denn wegen $A_n \in A_1$ ist $X \cap A_1$ nichtleer. Wegen $A_{i-1} \in A_i$ muss auch $Y \neq A_i$ für $1 < i \leq n$ sein. Es gibt daher kein derartiges Y – ein Widerspruch zum Fundierungsaxiom.

Somit kann es keine derartige verschachtelte Mengenkette geben. □

1.3 Konstruktion der natürlichen Zahlen

Die Axiome nach Zermelo und Fraenkel garantieren die Existenz bestimmter Mengen (wie beispielsweise das Nullmengenaxiom und das Unendlichkeitsaxiom), und erlauben Methoden, mit denen weitere Mengen konstruiert werden können. Wie aber entstehen aus Mengen die uns bekannten Zahlen Null, Eins, Zwei, Drei, ...?

Betrachten wir zunächst folgende einfache Rechnungen:

$$\begin{aligned}
 1 &= 0 + 1 \\
 2 &= 1 + 1 \\
 3 &= 2 + 1 \\
 4 &= 3 + 1 \\
 5 &= 4 + 1 \\
 &\vdots
 \end{aligned}
 \tag{1.1}$$

Das Beispiel zeigt naiv, dass sich jede natürliche Zahl n mit Ausnahme von Null als Summe einer anderen natürlichen Zahl plus Eins anschreiben lässt: $n = (n - 1) + 1$. Wir nennen n auch den *Nachfolger* von $n - 1$: Eins ist der Nachfolger von Null, Zwei ist der Nachfolger von Eins, und so weiter. Der Nachfolgerbegriff steht im Zentrum unserer Definition der natürlichen Zahlen.

Eine ähnliche Eigenschaft besitzt die Menge, deren Existenz uns das Unendlichkeitsaxiom garantiert. Diese enthält für jedes Element X die Menge $\{X, \{X\}\}$. Da zusätzlich auch die leere Menge $\emptyset$ enthalten ist, kann man – analog zum Beispiel mit den Zahlen – auch hier fortgesetzt beliebig viele Element der Menge anschreiben, indem man diese Eigenschaft auf $\emptyset$ anwendet, und dann sukzessive auf die erzeugten neuen Elemente:

$$\emptyset \rightarrow \{\emptyset\} \rightarrow \{\emptyset, \{\emptyset\}\} \rightarrow \{\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\}\} \rightarrow \dots
 \tag{1.2}$$

Es lässt sich also $\{\emptyset\}$ gewissermaßen als „Nachfolger“ von $\emptyset$ ansehen. Das wird durch folgende Definition formalisiert.

Definition 1.3.1 ([50, S. 288]). Sei X eine Menge. Dann ist der Nachfolger $S(X)$ von X definiert durch

$$S(X) := X \cup \{X\}.$$

Falls X eine Menge ist, so existiert $S(X)$ wegen Axiom 3 und Axiom 4. Die Eigenschaft, dass eine Menge M die leere Menge $\emptyset$ enthält und jedes Element $X \in M$ einen Nachfolger besitzt, wollen wir Nachfolgereigenschaft nennen.

Definition 1.3.2 ([50, S. 288]). Für eine Menge M definieren wir die Nachfolgereigenschaft $\delta(M)$ durch

$$\delta(M) := \emptyset \in M \wedge (\forall X \in M : S(X) \in M).$$

Vergleichen wir Beispiel 1.1 mit Beispiel 1.2, zeigen sich (gewollte) Parallelen. So gibt es in beiden Beispielen ein ausgezeichnetes Element, das am Anfang der Kette steht: In Beispiel 1.1 ist das Null, während diese Position in Beispiel 1.2 von der leeren Menge $\emptyset$ übernommen wird. In beiden gibt es zudem eine Vorschrift, wie aus einer gegebenen Menge ihr Nachfolger erzeugt werden kann. Wurde in 1.1 Eins als Nachfolger von Null angesehen, ist in 1.2 die Menge $\{\emptyset\}$ der Nachfolger von $\emptyset$. Es läge also nahe, die leere Menge als Null zu definieren, ihren Nachfolger als Eins, und so weiter.

Es reicht allerdings nicht aus, die Menge der natürlichen Zahlen einfach als jene Menge zu definieren, deren Existenz das Unendlichkeitsaxiom sicherstellt. Dieses Axiom garantiert lediglich die Existenz einer Menge, welche die Nachfolgereigenschaft besitzt – über die Eindeutigkeit oder eventuelle weitere Elemente der Menge trifft es keine Aussagen. Würden wir also einfach $\mathbb{N} := X$ setzen (mit X aus Axiom 6), wäre die Existenz „exotischer“ natürlicher Zahlen nicht auszuschließen.

Ziel ist also eine Menge der Gestalt $\{\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\}, \{\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\}\}, \dots\}$. Um diese zu erhalten, bilden wir den Durchschnitt über alle Teilmengen der Menge X aus Axiom 6, für welche die Nachfolgereigenschaft gilt. Das Resultat definieren wir als die Menge der natürlichen Zahlen $\mathbb{N}$.

Definition 1.3.3 ([50, S. 289]). Sei X die Menge aus Axiom 6. Dann ist die Menge der natürlichen Zahlen $\mathbb{N}$ definiert als

$$\mathbb{N} := \bigcap \{Y \in \mathcal{P}(X) \mid \delta(Y)\}.$$

Die Definition von $\mathbb{N}$ ist wegen Axiom 5 zulässig.

Proposition 1.3.4 ([50, S. 289]). Die Menge der natürlichen Zahlen $\mathbb{N}$ erfüllt die Nachfolgereigenschaft.

Beweis. Zur Vereinfachung setzen wir $\mathcal{N} := \{Y \in \mathcal{P}(X) \mid \delta(Y)\}$. Da für alle $Y \in \mathcal{N}$ die Nachfolgereigenschaft gilt, ist stets $\emptyset \in Y$, und damit auch $\emptyset \in \bigcap \mathcal{N}$, also $\emptyset \in \mathbb{N}$.

Sei nun $X \in \mathbb{N}$. Dann ist $X \in Y$ für alle $Y \in \mathcal{N}$. Da Y die Nachfolgereigenschaft erfüllt, müssen alle Y auch $S(X)$ enthalten. Damit ist $S(X) \in \bigcap \mathcal{N}$, also $S(X) \in \mathbb{N}$. $\square$

Proposition 1.3.5 ([50, S. 289]). Falls Z die Nachfolgereigenschaft erfüllt, so ist $\mathbb{N} \subseteq Z$. Die natürlichen Zahlen sind somit die kleinste Menge, für welche die Nachfolgereigenschaft gilt.

Beweis. Angenommen, die Menge Z erfüllt die Nachfolgereigenschaft. Sei X wie in Axiom 6 und $\mathcal{N}$ wie in Proposition 1.3.4. Dann ist jedenfalls $\emptyset \in (Z \cap X)$. Für ein $U \in (Z \cap X)$ ist $S(U) \in Z$ sowie $S(U) \in X$ da Z und X die Nachfolgereigenschaft besitzen. Damit ist $S(U) \in (Z \cap X)$. Es erfüllt also auch $Z \cap X$ die Nachfolgereigenschaft.

Außerdem ist $(Z \cap X) \subseteq X$, und da $Z \cap X$ die Nachfolgereigenschaft besitzt, gilt $(Z \cap X) \in \mathcal{N}$. Wegen $\mathbb{N} = \bigcap \mathcal{N}$ ist $\mathbb{N} \subseteq (Z \cap X)$.

Insgesamt ist daher $\mathbb{N} \subseteq (Z \cap X) \subseteq Z$, woraus $\mathbb{N} \subseteq Z$ folgt. $\square$

Die übliche Schreibweise der natürlichen Zahlen erhalten wir durch folgende Definitionen.

$$0 := \emptyset$$

$$1 := \{\emptyset\}$$

$$2 := \{\emptyset, \{\emptyset\}\}$$

$$\begin{aligned} 3 &:= \{\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\} \\ &\vdots \end{aligned}$$

Ein weiteres Axiomensystem, durch welches die natürlichen Zahlen beschrieben werden können, sind die Peano-Axiome. Diese wurden von Giuseppe Peano im Jahr 1889 vorgestellt, und charakterisieren die natürlichen Zahlen durch fünf Axiome [47]. In unserem Modell werden die Peano-Axiome zu beweisbaren Aussagen.

Satz 1.3.6 ([50, S. 290; 34, S. 2]). *Die Menge der natürlichen Zahlen $\mathbb{N}$ erfüllt die Peano-Axiome:*

- (1) *Null ist eine natürliche Zahl: $\emptyset \in \mathbb{N}$.*
- (2) *Jede natürliche Zahl hat einen natürlichen Nachfolger: $\forall n \in \mathbb{N} : (S(n) \in \mathbb{N})$.*
- (3) *Null ist Nachfolger keiner Zahl: $\forall n \in \mathbb{N} : (S(n) \neq \emptyset)$.*
- (4) *Jede natürliche Zahl ist Nachfolger höchstens einer Zahl:*

$$\forall n, m \in \mathbb{N} : ((S(n) = S(m)) \Rightarrow n = m).$$

- (5) *Enthält eine Teilmenge der natürlichen Zahlen die Null, und mit jeder Zahl ihren Nachfolger, so entspricht diese Menge den natürlichen Zahlen (Induktionsprinzip):*

$$\forall N \in \mathcal{P}(\mathbb{N}) : ((\emptyset \in N \wedge \forall n \in N : S(n) \in N) \Rightarrow N = \mathbb{N}).$$

Beweis. (1) und (2) folgen direkt aus Proposition 1.3.4.

Angenommen, es gäbe eine natürliche Zahl n mit $S(n) = \emptyset$. Dann wäre $n \in n \cup \{n\} = S(n) = \emptyset$, also $n \in \emptyset$ – Widerspruch. Somit kann es kein solches $n \in \mathbb{N}$ geben.

Zu (4): Seien $n, m \in \mathbb{N}$ mit $S(n) = S(m)$. Angenommen $m = \emptyset$, dann wäre $\{\emptyset\} = S(\emptyset) = S(m)$ und $m \in S(m)$, also $m \in \{\emptyset\}$, woraus $m = \emptyset$ folgt. Es ist also $m = n$. Analog zeigt sich im Fall $n = \emptyset$, dass $m = n$ ist.

Nehmen wir daher an, dass $m \neq \emptyset$ und $n \neq \emptyset$ sind. Sei $k \in n$, dann ist $k \in n \cup \{n\} = S(n) = S(m) = m \cup \{m\}$. Es ist demnach entweder $k \in m$ oder $k = m$.

Angenommen $k = m$. Dann ist $n \in n \cup \{n\} = S(n) = S(m) = S(k) = k \cup \{k\}$, also $n \in k$ oder $n = k$. Im Fall $n = k$ folgt nach der Definition von k , dass $n \in n$ ist, was nach Proposition 1.2.1 nicht sein kann. Auch der Fall $n \in k$ mit $k \in n$ führt deswegen auf einen Widerspruch.

Somit bleibt nur, dass $k \in m$ ist, und da k beliebig in n war, ist $n \subseteq m$. Durch Vertauschen von m und n im obigen Beweis lässt sich zeigen, dass $m \subseteq n$ ist, woraus $n = m$ folgt.

Zu (5): Da N die Nachfolgereigenschaft erfüllt, ist $\mathbb{N} \subseteq N$ wegen Proposition 1.3.5. Weiters gilt $N \subseteq \mathbb{N}$, woraus $N = \mathbb{N}$ folgt. $\square$

Das Induktionsprinzip verwenden wir künftig in der gewohnten Form, wobei wir die Menge N implizit verwenden, und die Nachfolgereigenschaft beweisen.

1.4 Rechenoperationen auf $\mathbb{N}$

Für das Addieren in den natürlichen Zahlen verwenden wir eine Idee, die vielen Schülern in der Volksschule geläufig ist: Lautet die Aufgabe zum Beispiel, $5 + 4$ zu berechnen, halten viele Kinder zunächst fünf Finger hoch, und strecken *der Reihe nach* vier weitere Finger aus, während sie dabei leise mitzählen. Sind auf der zweiten Hand dann vier Finger ausgestreckt, ist das Ergebnis der Addition die zuletzt leise mitgesprochene Zahl.

Der Vorgang „Finger der Reihe nach ausstrecken“ lässt sich in unserem Modell dadurch ausdrücken, sukzessive die Nachfolger einer Zahl zu bestimmen – und zwar so oft, wie es der zweite Summand in der Addition angibt. Ziel ist es also, die Rechnung $5 + 4$ auf $(S \circ S \circ S \circ S)(5)$ zurückzuführen, wo vier Mal die Funktion S hintereinander auf 5 angewendet wird. Das gelingt, indem wir einen rekursiven Ansatz für die Definition wählen. Wir legen zunächst fest, dass $n + S(m) = S(n + m)$ sein soll. Dadurch, dass sich $n + m$ für $m \neq 0$ wieder als $n + S(m')$ anschreiben lässt (mit $S(m') = m$), kann man erneut in die Definition einsetzen, und erhält so eine Kette von Funktionen. Damit diese Kette ein Ende findet, legen wir zudem fest, dass $n + 0 = n$ sein soll.

Formal ausgedrückt verwenden wir für die Definition des Additionsoperators eine zweistellige Funktion $+$: $\mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$, wobei wir $+(n, m)$ in der üblichen Infixnotation als $n + m$ anschreiben.

Definition 1.4.1 ([50, S. 293]). Seien $n, m \in \mathbb{N}$. Die Addition auf den natürlichen Zahlen ist definiert durch $+$: $\mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$, mit

$$\begin{aligned} n + 0 &:= n \\ n + S(m) &:= S(n + m). \end{aligned}$$

Aus der obigen „Definition“ wird aber nicht klar, ob es tatsächlich eine solche Funktion $+$: $\mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ gibt, oder sie eindeutig ist – denn schließlich wurden nur wünschenswerte Eigenschaften dieser Funktion angeschrieben. Der Dedekindsche Rekursionssatz zeigt, dass diese und ähnliche rekursive Definitionen möglich sind.

Satz 1.4.2 ([34, S. 8–10]). Sei A eine Menge, $a \in A$ und $g : A \rightarrow A$ eine Funktion. Dann gibt es genau eine Abbildung $f : \mathbb{N} \rightarrow A$ mit $f(0) = a$ und $f(S(n)) = g(f(n))$ für alle $n \in \mathbb{N}$.

Beweis. Wir führen den Beweis in zwei Schritten. Zunächst zeigen wir, dass eine solche Funktion f , wenn sie existiert, eindeutig ist. Anschließend zeigen wir die Existenz dieser Funktion.

Zur Eindeutigkeit: Angenommen, es gäbe zwei Funktionen $f_1 : \mathbb{N} \rightarrow A$ und $f_2 : \mathbb{N} \rightarrow A$ mit den geforderten Eigenschaften. Sei $N := \{n \in \mathbb{N} \mid f_1(n) = f_2(n)\}$, also die Menge aller natürlichen Zahlen, für welche f_1 und f_2 die selben Werte haben. Dann ist $f_1(0) = a = f_2(0)$, und somit $0 \in N$.

Sei nun $n \in N$, dann ist

$$f_1(S(n)) = g(f_1(n)) = g(f_2(n)) = f_2(S(n)),$$

und somit auch $S(n) \in N$. Aufgrund des Induktionsprinzips folgt, dass $N = \mathbb{N}$ sein muss. Es ist daher $f_1 = f_2$.

Für den Beweis der Existenz der Funktion f arbeiten wir mit der Definition einer Funktion als Teilmenge des kartesischen Produkts von Definitions- und Wertebereich:

$$f = (\mathbb{N}, A, G) \text{ mit } G \subseteq \mathbb{N} \times A.$$

Wir betrachten die Mengenfamilie $\mathcal{G} \subseteq \mathcal{P}(\mathbb{N} \times A)$ mit den folgenden Eigenschaften:

- (1) $\forall G \in \mathcal{G}: (0, a) \in G$
- (2) $\forall G \in \mathcal{G}: (n, b) \in G \Rightarrow (S(n), g(b)) \in G.$

Ein Beispiel für eine Menge mit diesen Eigenschaften wäre $\mathbb{N} \times A$: Diese enthält jedenfalls $(0, a)$, da $a \in A$ ist, und für ein (n, b) in dieser Menge ist auch $(S(n), g(b))$ enthalten, da die Funktion g von A nach A abbildet. Nicht erfüllt ist hingegen die *Rechtseindeutigkeit*. Damit es sich bei dem Tupel $G' = \mathbb{N} \times A$ um eine Funktion handelt, müsste gelten, dass jedem Element aus $\mathbb{N}$ *genau ein* Element aus A zugeordnet wird. Das ist hier nicht notwendigerweise erfüllt.

Wir gehen daher ähnlich vor wie bei der Konstruktion der natürlichen Zahlen, und bilden den Durchschnitt über die Menge $\mathcal{G}$. Anschließend zeigen wir über Induktion, dass es sich bei dieser Menge um eine Funktion handelt.

Sei also $G := \bigcap \mathcal{G}$ und $N := \{n \in \mathbb{N} \mid \exists! b \in A : (n, b) \in G\}$. Nachdem $(0, a)$ in allen Elementen von $\mathcal{G}$ enthalten ist, muss es auch im Durchschnitt, also G enthalten sein. Da es sich bei G um den Durchschnitt aller Mengen mit der Eigenschaft (2) handelt, erfüllt auch G diese Eigenschaft.

Angenommen, es wäre $(0, c) \in G$ mit $c \neq a$. Dann ist die Menge $G' := G \setminus \{(0, c)\}$ eine echte Teilmenge von G , für die ebenfalls die Eigenschaften (1) und (2) gelten (das folgt daher, dass diese Eigenschaften wie zuvor gezeigt für G gelten, und wir aus G' nur ein Element entfernen, was keine Auswirkungen auf die Implikation in (2) hat). Es ist also $G' \in \mathcal{G}$, und damit $G \subseteq G'$, also $(0, c) \in G'$ – ein Widerspruch. Somit ist $0 \in N$.

Sei nun $n \in N$. Per Definition von N gibt es genau ein $b \in A$, sodass $(n, b) \in G$ ist. Aufgrund von Eigenschaft (2) folgt, dass auch $(S(n), g(b)) \in G$ ist. Für den Beweis der Eindeutigkeit nehmen wir wieder an, dass es ein $c \neq g(b)$ gäbe, sodass $(S(n), c) \in G$. Dann ist $G' := G \setminus \{(S(n), c)\}$ eine echte Teilmenge von G , und mit den selben

Argumenten wie zuvor folgt, dass auch G' die Eigenschaften (1) und (2) erfüllt. Wir erhalten daher wieder $G \subseteq G'$, also $(S(n), c) \in G'$ – ein Widerspruch. Damit ist auch die Eindeutigkeit für $S(n)$ gezeigt, und $S(n) \in N$. Aufgrund des Induktionsprinzips folgt, dass $N = \mathbb{N}$ ist.

Wir haben daher gezeigt, dass $G \subseteq \mathbb{N} \times A$ die Rechtseindeutigkeit erfüllt. Somit ist $f = (\mathbb{N}, A, G)$ die gesuchte Abbildung. $\square$

Wir können nun über Rekursionssatz beweisen, dass die Addition eindeutig definiert ist.

Proposition 1.4.3 ([20, S. 16; 18, S. 42]). *Es gibt genau eine Funktion $+$: $\mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ mit den Eigenschaften aus Definition 1.4.1.*

Beweis. Sei $n \in \mathbb{N}$ beliebig, aber fest. Setzen wir $g := S$, $A := \mathbb{N}$ und $a := n$, so folgt aus dem Rekursionssatz, dass es genau eine Funktion $f_n : \mathbb{N} \rightarrow \mathbb{N}$ gibt, für die

$$\begin{aligned} f_n(0) &:= n, \\ f_n(S(m)) &:= S(f_n(m)) \end{aligned}$$

erfüllt ist. Definieren wir $+(n, m) := f_n(m)$, so ist $+(n, 0) = f_n(0) = n$ und

$$+(n, S(m)) = f_n(S(m)) = S(f_n(m)) = S(+(n, m)),$$

also genau die Eigenschaften aus Definition 1.4.1. $\square$

Um zu zeigen, dass es sich bei den natürlichen Zahlen mit der Addition um ein kommutatives Monoid handelt, benötigen wir folgenden Hilfssatz.

Lemma 1.4.4 ([50, S. 295; 34, S. 4]). *Seien $n, m \in \mathbb{N}$. Falls 0 das neutrale Element der Addition ist, so gilt $n + S(m) = S(n) + m$.*

Beweis. Sei $n \in \mathbb{N}$ beliebig, aber fest. Dann ist

$$n + S(0) = S(n + 0) = S(0 + n) = 0 + S(n),$$

und daher der Induktionsanfang erfüllt. Angenommen, für $m \in \mathbb{N}$ gilt die Aussage. Dann ist

$$n + S(S(m)) = S(n + S(m)) = S(S(n) + m) = S(n) + S(m),$$

also ist die Aussage auch für $S(m)$ wahr, und wegen des Induktionsprinzips daher für alle $m \in \mathbb{N}$. Da n beliebig war, gilt die Aussage somit für alle $n, m \in \mathbb{N}$. $\square$

Proposition 1.4.5 ([50, S. 295]). $(\mathbb{N}, +, 0)$ ist ein kommutatives Monoid. Es sind somit folgende Eigenschaften erfüllt:

(1) $\forall k, l, m \in \mathbb{N} : (k + l) + m = k + (l + m)$ (Assoziativität).

(2) $\forall n \in \mathbb{N} : n + 0 = 0 + n = n$ (Nullelement).

(3) $\forall n, m \in \mathbb{N} : n + m = m + n$ (Kommutativität).

Beweis. Zu (1): Seien $k, l \in \mathbb{N}$ beliebig, aber fest. Dann ist $(k + l) + 0 = k + l = k + (l + 0)$. Angenommen, es ist $(k + l) + m = k + (l + m)$ für ein $m \in \mathbb{N}$ erfüllt. Dann ist

$$(k + l) + S(m) = S((k + l) + m) = S(k + (l + m)) = k + S(l + m) = k + (l + S(m)).$$

Wegen des Induktionsprinzips und der Beliebigkeit von k und l folgt damit, dass $(k + l) + m = k + (l + m)$ für alle $k, l, m \in \mathbb{N}$ erfüllt ist.

Zu (2): Aus Definition 1.4.1 wissen wir, dass 0 rechtsneutral für alle $n \in \mathbb{N}$ ist. Wir zeigen über vollständige Induktion, dass 0 auch für alle natürlichen Zahlen linksneutral ist. Der Induktionsanfang folgt direkt aus der Definition, es ist $0 + 0 = 0$. Ist $0 + n = n$ für ein $n \in \mathbb{N}$, so ist $0 + S(n) = S(0 + n) = S(n)$, also ist 0 auch linksneutral für $S(n)$, und damit wegen des Induktionsprinzips für alle natürlichen Zahlen.

Zum Beweis von (3): Sei $n \in \mathbb{N}$ beliebig, aber fest. Dann ist wegen (2) $n + 0 = 0 + n$. Gilt $n + m = m + n$ für ein $m \in \mathbb{N}$, dann ist

$$n + S(m) = S(n + m) = S(m + n) = m + S(n) = S(m) + n,$$

wobei für den letzten Umformungsschritt Lemma 1.4.4 verwendet wurde. Nach dem Induktionsprinzip ist daher $n + m = m + n$ für alle $m \in \mathbb{N}$, und da n beliebig war, auch für alle $n \in \mathbb{N}$. $\square$

Aufgrund der Assoziativität der Addition vereinbaren wir, dass wir $(a + b) + c$ auch ohne Klammern als $a + b + c$ schreiben.

Wir sind nun dazu in der Lage, die Identität $1 + 1 = 2$ zu beweisen. Dabei verwenden wir $1 := S(0)$, und $2 := S(1)$.

$$\begin{aligned}
 1 + 1 &= 1 + S(0) && \text{(Definition von 1)} \\
 &= S(1 + 0) && \text{(Definition von +)} \\
 &= S(1) && \text{(Definition von +)} \\
 &= 2 && \text{(Definition von 2)}
 \end{aligned}$$

Im obigen Beispiel haben wir Zwei, also den Nachfolger von Eins, dadurch erhalten, dass wir zu Eins die Zahl Eins addiert haben. Das Prinzip lässt sich auch verallgemeinern, wie folgende Proposition zeigt.

Proposition 1.4.6. *Sei $n \in \mathbb{N}$. Dann ist $n + 1 = 1 + n = S(n)$.*

Beweis. Per Definition der Addition ist $n + 1 = n + S(0) = S(n + 0) = S(n)$, und da die Addition kommutativ ist, gilt $1 + n = n + 1 = S(n)$. $\square$

Die Definition der Multiplikation gelingt, ähnlich wie bei der Addition, rekursiv. Wir fordern zunächst, dass die Multiplikation einer natürlichen Zahl mit Null das Ergebnis Null hat. Für die Multiplikation einer Zahl ungleich Null verwenden wir die Idee der Multiplikation als „verkürzte Addition“:

$$n \cdot m = n \cdot (m - 1) + n = n \cdot (m - 2) + n + n = \dots = \underbrace{n + n + n + \dots + n}_{m \text{ mal}}$$

Definition 1.4.7 ([50, S. 294]). Seien $n, m \in \mathbb{N}$. Die Multiplikation $\cdot : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ auf den natürlichen Zahlen ist definiert durch

$$\begin{aligned} n \cdot 0 &:= 0, \\ n \cdot S(m) &:= (n \cdot m) + n. \end{aligned}$$

Da es sich auch hier um eine rekursive Definition handelt, ist, wie bei der Addition, zu zeigen, dass es tatsächlich eine solche Funktion gibt.

Proposition 1.4.8. *Es gibt genau eine Funktion $\cdot : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ mit den Eigenschaften aus Definition 1.4.7.*

Beweis. Sei $n \in \mathbb{N}$ beliebig, aber fest. Nach Satz 1.4.2 gibt es genau eine Funktion $f_n : \mathbb{N} \rightarrow \mathbb{N}$ mit

$$\begin{aligned} f_n(0) &:= 0, \\ f_n(S(m)) &:= f_n(m) + n. \end{aligned}$$

Setzen wir $\cdot(n, m) := f_n(m)$, so ist $\cdot(n, 0) = f_n(0) = 0$ und

$$\cdot(n, S(m)) = f_n(S(m)) = f_n(m) + n = \cdot(n, m) + n,$$

also genau die Eigenschaften aus Definition 1.4.7. □

Damit wir Aussagen über die algebraische Struktur von $(\mathbb{N}, \cdot)$ treffen können, beweisen wir folgende Hilfssätze.

Lemma 1.4.9 ([50, S. 295]). *Sei $n \in \mathbb{N}$. Dann ist $0 \cdot n = 0$.*

Beweis. Wir führen den Beweis über vollständige Induktion. Es ist $0 \cdot 0 = 0$ wegen der Definition der Multiplikation, und somit der Induktionsanfang erfüllt. Angenommen, für ein $n \in \mathbb{N}$ ist $0 \cdot n = 0$. Dann ist $0 \cdot S(n) = (0 \cdot n) + 0 = 0 + 0 = 0$. Wegen des Induktionsprinzips folgt somit, dass $0 \cdot n = 0$ für alle $n \in \mathbb{N}$ ist. □

Lemma 1.4.10 ([50, S. 296]). *Seien $n, m \in \mathbb{N}$. Dann ist $S(n) \cdot m = (n \cdot m) + m$.*

Beweis. Wir verwenden erneut das Induktionsprinzip für den Beweis. Sei $n \in \mathbb{N}$ beliebig, aber fest. Es ist dann $S(n) \cdot 0 = 0 = (n \cdot 0) + 0$, und somit der Induktionsanfang erfüllt. Nehmen wir nun an, die Aussage gilt für ein $m \in \mathbb{N}$. Dann ist

$$\begin{aligned} S(n) \cdot S(m) &= (S(n) \cdot m) + S(n) = \\ &= ((n \cdot m) + m) + S(n) = \\ &= (n \cdot m) + S(n) + m = \\ &= ((n \cdot m) + n) + S(m) = \\ &= (n \cdot S(m)) + S(m), \end{aligned}$$

wobei wir für die Umformungen die Assoziativität und Kommutativität der Addition, sowie Lemma 1.4.4 verwendet haben. Aufgrund des Induktionsprinzips und der Beliebigkeit von $n \in \mathbb{N}$ folgt, dass die Aussage für alle $n, m \in \mathbb{N}$ erfüllt ist. $\square$

Bemerkung. Bei der Aussage dieses Lemmas handelt es sich um eine abgeschwächte Form des Distributivgesetzes: $(n + 1) \cdot m = (n \cdot m) + m$.

Proposition 1.4.11 ([50, S. 295–296]). $(\mathbb{N}, +, \cdot, 0, 1)$ ist ein kommutativer Halbring mit Eins. Es sind daher folgende Eigenschaften erfüllt.

- (1) $(\mathbb{N}, +, 0)$ ist ein kommutatives Monoid.
- (2) $\forall k, l, m \in \mathbb{N} : (k \cdot l) \cdot m = k \cdot (l \cdot m)$ (Assoziativität).
- (3) $\forall n \in \mathbb{N} : n \cdot 1 = 1 \cdot n = n$ (Einselement).
- (4) $\forall n, m \in \mathbb{N} : n \cdot m = m \cdot n$ (Kommutativität).
- (5) $\forall k, l, m \in \mathbb{N} : k \cdot (l + m) = (k \cdot l) + (k \cdot m)$ (Distributivität).

Beweis. Die erste Aussage wurde in Proposition 1.4.5 gezeigt.

Zur Kommutativität: Sei $n \in \mathbb{N}$ beliebig, aber fest. Dann ist $n \cdot 0 = 0 = 0 \cdot n$ wegen Lemma 1.4.9. Angenommen, für $m \in \mathbb{N}$ wäre $n \cdot m = m \cdot n$. Dann ist

$$n \cdot S(m) = (n \cdot m) + n = (m \cdot n) + n = S(m) \cdot n,$$

wobei für den letzten Umformungsschritt Lemma 1.4.10 verwendet wurde. Aufgrund des Induktionsprinzips und der Beliebigkeit von n folgt, dass die Multiplikation kommutativ ist.

Zum Einselement: Sei $n \in \mathbb{N}$. Dann ist $n \cdot 1 = n \cdot S(0) = (n \cdot 0) + n = n$, also ist 1 rechtsneutral für $n \in \mathbb{N}$, und wegen der Kommutativität der Multiplikation auch linksneutral.

Zur Distributivität: Seien $l, m \in \mathbb{N}$ beliebig, aber fest. Dann ist $0 \cdot (l + m) = 0 = 0 \cdot l + 0 \cdot m$. Angenommen, die Aussage wäre für ein $k \in \mathbb{N}$ erfüllt. Dann ist wegen Lemma 1.4.10

$$\begin{aligned} S(k) \cdot (l + m) &= (k \cdot (l + m)) + (l + m) = \\ &= (k \cdot l) + (k \cdot m) + (l + m) = \\ &= ((k \cdot l) + l) + ((k \cdot m) + m) = \\ &= (S(k) \cdot l) + (S(k) \cdot m). \end{aligned}$$

Wegen des Induktionsprinzips, der Beliebigkeit von l und m sowie der Kommutativität der Multiplikation folgt, dass die Distributivgesetze erfüllt sind.

Zur Assoziativität: Seien $l, m \in \mathbb{N}$ beliebig, aber fest. Dann ist

$$(0 \cdot l) \cdot m = 0 \cdot m = 0 = 0 \cdot (l \cdot m).$$

Angenommen, die Aussage wäre für ein $k \in \mathbb{N}$ erfüllt. Es ist dann

$$\begin{aligned} (S(k) \cdot l) \cdot m &= ((k \cdot l) + l) \cdot m = \\ &= ((k \cdot l) \cdot m) + (l \cdot m) = \\ &= (k \cdot (l \cdot m)) + (l \cdot m) = \\ &= (k + 1) \cdot (l \cdot m) = \\ &= S(k) \cdot (l \cdot m), \end{aligned}$$

wobei wir insbesondere Lemma 1.4.10, die Existenz des Einselements sowie die Distributivität verwendet haben. Wegen des Induktionsprinzips sowie der Beliebigkeit von l und m folgt die Assoziativität. $\square$

Wie bereits zuvor bei der Addition, vereinbaren wir auch bei der Multiplikation, dass wir aufgrund der Assoziativität für $k, l, m \in \mathbb{N}$ den Ausdruck $k \cdot (l \cdot m)$ auch ohne Klammern als $k \cdot l \cdot m$ anschreiben dürfen.

Zudem legen wir die aus der Schule bekannte „Punkt vor Strich“-Regel fest: Für $k, l, m \in \mathbb{N}$ schreiben wir $(k \cdot l) + m$ als $k \cdot l + m$ an; wir sagen dazu, dass die Multiplikation „stärker“ als die Addition ist. Es ist – gerade auch für die Einführung dieser Regel in der Schule – erwähnenswert, dass es sich dabei nicht um ein beweisbares Rechengesetz, sondern um eine Konvention handelt.

Als Beispiel berechnen wir $2 \cdot 3$:

$$\begin{aligned} 2 \cdot 3 &= 2 \cdot 2 + 2 \\ &= 2 \cdot 1 + 2 + 2 \\ &= 2 \cdot 0 + 2 + 2 + 2 \\ &= 0 + 2 + 2 + 2 \\ &= 6 \end{aligned}$$

Aus Satz 1.3.6 ist bekannt, dass jede natürliche Zahl einen Nachfolger hat. Die folgende Proposition zeigt, dass jede natürliche Zahl außer Null auch einen Vorgänger besitzt.

Proposition 1.4.12. *Sei $n \in \mathbb{N}$. Dann ist $n = 0$, oder aber es gibt genau ein $n' \in \mathbb{N}$ mit $S(n') = n$.*

Beweis. Wir verwenden wieder das Induktionsprinzip. Für $n = 0$ ist die Aussage erfüllt. Angenommen, sie wäre es für ein $n \in \mathbb{N}$. Dann ist klarerweise $S(n) = S(n)$, also n der Vorgänger von $S(n)$, der wegen Satz 1.3.6 eindeutig ist. Die Aussage gilt somit wegen des Induktionsprinzips für alle $n \in \mathbb{N}$. $\square$

Wir können nun beweisen, dass die natürlichen Zahlen nullteilerfrei sind: Ist das Produkt zweier Zahlen gleich Null, dann ist es auch zumindest einer der Faktoren.

Proposition 1.4.13 ([50, S. 298]). *Die natürlichen Zahlen sind nullteilerfrei: Ist $n \cdot m = 0$, dann folgt $n = 0$ oder $m = 0$.*

Beweis. Seien $n, m \in \mathbb{N}$ mit $n \cdot m = 0$. Ist $n = 0$ oder $m = 0$, dann sind wir fertig. Andernfalls gibt es nach Proposition 1.4.12 $n', m' \in \mathbb{N}$ mit $S(n') = n$ und $S(m') = m$. Es ist dann

$$\begin{aligned} n \cdot m &= S(n') \cdot S(m') = \\ &= (n' + 1) \cdot (m' + 1) = \\ &= n' \cdot m' + n' + m' + 1 = \\ &= S(n' \cdot m' + n' + m'). \end{aligned}$$

Da der Nachfolger einer natürlichen Zahl nach Satz 1.3.6 immer ungleich Null ist, muss daher auch $n \cdot m \neq 0$ sein. $\square$

Ähnlich zur Nullteilerfreiheit gilt: Falls die Summe zweier natürlicher Zahlen gleich Null ist, dann sind es bereits beide Summanden. Daraus folgt insbesondere, dass es sich bei den natürlichen Zahlen mit der Addition um keine Gruppe handelt, da außer Null kein Element ein additives Inverses besitzt.

Proposition 1.4.14. *Seien $n, m \in \mathbb{N}$ mit $n + m = 0$. Dann ist $n = m = 0$.*

Beweis. Angenommen, es wäre $n \neq 0$ oder $m \neq 0$. Ohne Beschränkung der Allgemeinheit nehmen wir an, es wäre $n \neq 0$. Dann gibt es nach Proposition 1.4.12 ein $n' \in \mathbb{N}$ mit $S(n') = n$.

Es ist dann $n + m = S(n') + m = S(n' + m) = 0$ – ein Widerspruch zu Satz 1.3.6, da Null Nachfolger keiner natürlichen Zahl ist. Somit muss $n = 0$ und $m = 0$ gelten. $\square$

Für das Addieren in den natürlichen Zahlen gilt folgende Kürzungsregel.

Proposition 1.4.15 ([50, S. 298–299]). *Seien $n, m, k \in \mathbb{N}$. Dann folgt aus $n + k = m + k$, dass $n = m$ ist.*

Beweis. Wir verwenden erneut das Induktionsprinzip. Seien $n, m \in \mathbb{N}$ beliebig, aber fest. Dann folgt aus $n + 0 = m + 0$ direkt $n = m$. Nehmen wir nun an, für ein $k \in \mathbb{N}$ wäre $n + k = m + k$. Ist $n + S(k) = m + S(k)$, so folgt daraus wegen der Definition der Addition $S(n + k) = S(m + k)$, und hieraus wegen Satz 1.3.6 $n + k = m + k$. Anwenden

der Induktionsvoraussetzung liefert $n = m$. Wegen des Induktionsprinzips und der Beliebigkeit von n und m gilt die Aussage somit für alle $n, m, k \in \mathbb{N}$. $\square$

Eine ähnliche Rechenregel gilt auch für die Multiplikation (siehe Proposition 1.5.6). Für deren Beweis verwenden wir Eigenschaften der natürlichen Ordnung, die im nächsten Abschnitt vorgestellt wird.

1.5 Die natürliche Ordnung

In Abschnitt 1.1 wurde herausgearbeitet, dass der Mensch wahrscheinlich mit dem Zählen begonnen hat, um feststellen zu können, ob Anzahlen gleich oder unterschiedlich sind. Falls die Anzahlen verschieden sind, stellt sich die Frage, ob die eine Anzahl größer ist als die andere – denn von der Antwort auf diese Frage hängt es ab, ob eine Herde, die ein steinzeitlicher Mensch „gezählt“ hat, vollständig ist oder nicht [31, S. 28].

Greifen wir nochmals das Beispiel mit dem Hirten und der Schafherde von Seite 6 auf. Für jedes Schaf, das in der Früh an ihm vorbeigeht, würde er einen Kieselstein auf einen Haufen legen. Am Abend nimmt er für jedes Schaf, das wieder zurück in die Höhle geht, einen Kieselstein vom Haufen weg. Ist mit dem letzten Stein auch gleichzeitig das letzte Schaf am Hirten vorbeigegangen, so ist die Herde vollständig – soweit die in Abschnitt 1.1 vorgestellte Idee des konkreten Zählens mittels Hilfsmengen.

Für dieses Kapitel interessant sind die Fälle, in denen entweder Kieselsteine übrigbleiben, oder zu wenige vorhanden sind. Angenommen, es gäbe keine Steine mehr, aber es stehen weiterhin Schafe vor der Höhle. Dann erscheint es uns intuitiv als klar, dass zumindest bereits so viele Schafe in der Höhle sind, wie es in der Früh waren – denn die Größe des Haufens hat sich nicht verändert. Da aber weiterhin Schafe vor der Höhle stehen, muss die Herde insgesamt über den Tag größer geworden sein. Es sind also mehr Schafe vorhanden als zuvor.

Würde man die neue Anzahl der Schafe „zählen“ wollen, müsste man auf den Haufen weitere Kieselsteine hinzulegen. Anders ausgedrückt: Es gibt eine Zahl, die wir zur ursprünglichen Anzahl dazugeben können, um die neue Anzahl zu erhalten. Diese Idee wird durch folgende Definition präzisiert.

Definition 1.5.1 ([34, S. 6]). Seien $n, m \in \mathbb{N}$. Dann definieren wir

$$n \leq m :\Leftrightarrow \exists k \in \mathbb{N} : n + k = m.$$

Wir schreiben weiters $m \geq n$ für $n \leq m$. Ist $n \neq m$ und $n \leq m$, so schreiben wir $n < m$, beziehungsweise $m > n$.

Aufgrund von Proposition 1.4.15 ist in dieser Definition die Zahl k , falls sie existiert, eindeutig bestimmt. Falls $n < m$, muss dieses $k > 0$ sein.

Falls $n < m$ für $n, m \in \mathbb{N}$ ist, dann folgt daraus $n \leq m$.

Betrachten wir als Beispiel die Zahlen 1 und 3. Wir können $3 = 1 + 2$ anschreiben, was nach obiger Definition $1 \leq 3$ bedeutet. Genauer gilt wegen $2 \neq 0$ sogar $1 < 3$.

Wir können nun zeigen, dass es sich bei $\leq$ um eine Totalordnung handelt. Für zwei natürliche Zahlen ist daher stets $n \leq m$ oder $m \leq n$.

Proposition 1.5.2 ([34, S. 6–7]). Die Relation $\leq$ ist eine Totalordnung auf $\mathbb{N}$. Es gilt daher $n \leq m$ oder $m \leq n$ für alle $n, m \in \mathbb{N}$. Zudem erfüllt $\leq$ die Eigenschaften einer Ordnungsrelation:

- (1) $\forall n \in \mathbb{N} : n \leq n$ (Reflexivität).
- (2) $\forall n, m \in \mathbb{N} : (n \leq m \wedge m \leq n) \Rightarrow n = m$ (Antisymmetrie).
- (3) $\forall k, m, n \in \mathbb{N} : (k \leq m \wedge m \leq n) \Rightarrow k \leq n$ (Transitivität).

Beweis. Zur Reflexivität: Es ist $n + 0 = n$ für alle $n \in \mathbb{N}$, und daher nach Definition $n \leq n$.

Zur Antisymmetrie: Angenommen, für $n, m \in \mathbb{N}$ ist $n \leq m$ sowie $m \leq n$. Nach Definition 1.5.1 bedeutet das, dass es $k, k' \in \mathbb{N}$ mit $m = n + k$ sowie $n = m + k'$ gibt. Einsetzen der ersten Gleichung in die zweite liefert $n = n + (k + k')$. Wegen Proposition 1.4.15 ist $k + k' = 0$, woraus durch Proposition 1.4.14 $k = k' = 0$ folgt. Es ist daher $n = m$.

Zur Transitivität: Seien $k, m, n \in \mathbb{N}$ mit $k \leq m$ sowie $m \leq n$. Dann gibt es $l, l' \in \mathbb{N}$ mit $m = k + l$ und $n = m + l'$. Einsetzen der ersten Gleichung in die zweite liefert $n = (k + l) + l' = k + (l + l')$. Nach Definition 1.5.1 ist daher $k \leq n$.

Um zu beweisen, dass es sich bei $\leq$ um eine Totalordnung handelt, verwenden wir das Induktionsprinzip. Sei $n \in \mathbb{N}$ beliebig, aber fest. Dann behaupten wir, dass $n \leq m$ oder $m \leq n$ für alle $m \in \mathbb{N}$ erfüllt ist.

Im Fall $m = 0$ ist $0 \leq n$, da $n = 0 + n$ ist. Nehmen wir nun an, dass $n \leq m$ oder $m \leq n$ für ein $m \in \mathbb{N}$ erfüllt ist. Wir können die Fälle $m = n$, $n < m$ und $m < n$ unterscheiden, die sich gegenseitig ausschließen. Ist $m = n$, so ist $n \leq n + 1 = m + 1 = S(m)$, also $n \leq S(m)$.

Falls $m \neq n$ ist, bleiben nur die Fälle $m < n$ oder $n < m$. Angenommen, es ist $m < n$. Dann gibt es ein $k \in \mathbb{N} \setminus \{0\}$, sodass $n = m + k$ ist. Wegen $k \neq 0$ gibt es ein $k' \in \mathbb{N}$ mit $k = S(k')$. Es ist daher $n = m + S(k') = S(m) + k'$, und daher $S(m) \leq n$.

Ist $n < m$, so gibt es wieder ein $k \in \mathbb{N} \setminus \{0\}$ mit $m = n + k$. Es ist dann $S(m) = m + 1 = n + (k + 1) = n + S(k)$, und somit $n \leq S(m)$.

Wir erhalten somit, dass $n \leq S(m)$ oder $S(m) \leq n$ gilt. Wegen des Induktionsprinzips und der Beliebigkeit von $n \in \mathbb{N}$ folgt somit die Aussage, dass stets zwei natürliche Zahlen über $\leq$ vergleichbar sind. $\square$

Eine besondere Eigenschaft der natürlichen Zahlen, auf die wir später zurückkommen werden, ist, dass jede nichtleere Teilmenge von $\mathbb{N}$ ein kleinstes Element besitzt. Insbesondere haben die natürlichen Zahlen mit Null ein Minimum.

Proposition 1.5.3. Sei $n \in \mathbb{N}$. Dann ist $0 \leq n$.

Beweis. Es ist $n = 0 + n$, und daher $0 \leq n$. $\square$

Die natürliche Ordnung erfüllt die Ordnungsaxiome, wie folgende Proposition zeigt.

Proposition 1.5.4 ([50, S. 297]). Seien $k, m, n \in \mathbb{N}$. Dann gilt:

- (1) Es ist $n \leq m \Leftrightarrow n + k \leq m + k$.
- (2) Falls $n > 0$ und $m > 0$, dann ist $n \cdot m > 0$.

Beweis. Zu (1): Angenommen $n \leq m$, dann gibt es ein $l \in \mathbb{N}$ mit $m = n + l$. Addieren von k auf beiden Seiten der Gleichung liefert $m + k = (n + k) + l$, also $n + k \leq m + k$.

Sei nun $n + k \leq m + k$, dann gibt es wieder ein $l \in \mathbb{N}$, sodass $m + k = (n + k) + l$ ist. Wegen Proposition 1.4.15 gilt $m = n + l$, also $n \leq m$.

Zu (2): Angenommen $n > 0$ und $m > 0$. Dann ist insbesondere $n \neq 0$ sowie $m \neq 0$, und aus Proposition 1.4.13 folgt, dass $n \cdot m \neq 0$ gelten muss. Zusammen mit Proposition 1.5.3 ergibt sich daher $n \cdot m > 0$. $\square$

Die folgende Proposition rechtfertigt das Rechnen mit Ungleichungen.

Proposition 1.5.5 ([34, S. 8; 50, S. 297–298]). *Seien $k, l, m, n \in \mathbb{N}$. Dann gilt:*

- (1) *Falls $k \leq l$ und $m \leq n$, dann ist $k + m \leq l + n$ und $k \cdot m \leq l \cdot n$.*
- (2) *Falls $l \neq 0$ und $l \cdot m \leq l \cdot n$, dann ist $m \leq n$.*

Beweis. Zu (1): Angenommen $k \leq l$ und $m \leq n$. Dann gibt es $c, c' \in \mathbb{N}$ mit $l = k + c$ und $n = m + c'$. Durch Addieren der Gleichungen erhalten wir $l + n = k + m + (c + c')$, also $k + m \leq l + n$. Multiplizieren wir die Gleichungen, so erhalten wir durch Anwenden des Distributivgesetzes $l \cdot n = (k + c) \cdot (m + c') = k \cdot m + (k \cdot c' + c \cdot m + c \cdot c')$, also $k \cdot m \leq l \cdot n$.

Zu (2): Falls $l \neq 0$ und $l \cdot m \leq l \cdot n$ ist, dann gibt es ein $k \in \mathbb{N}$ mit $l \cdot n = l \cdot m + k$. Angenommen, die Behauptung wäre falsch, dann muss wegen Proposition 1.5.2 $m > n$ gelten. Es gibt daher ein $k' \in \mathbb{N} \setminus \{0\}$ mit $m = n + k'$, also $l \cdot m = l \cdot n + l \cdot k'$.

Setzen wir das in die erste Gleichung ein, erhalten wir $l \cdot n = l \cdot n + l \cdot k' + k$. Wegen der Kürzungsregel ist $0 = l \cdot k' + k$, und aus Proposition 1.4.14 folgt $l \cdot k' = 0$ sowie $k = 0$. Da es in den natürlichen Zahlen keine Nullteiler gibt, muss $k' = 0$ sein – ein Widerspruch zu $k' \in \mathbb{N} \setminus \{0\}$. Es kann somit nur $m \leq n$ sein. $\square$

Die Aussagen der beiden vorigen Propositionen sowie die Transitivität gelten entsprechend auch für $<$.

Wir können nun die Kürzungsregel für die Multiplikation in den natürlichen Zahlen beweisen (siehe Proposition 1.4.15).

Proposition 1.5.6 ([50, S. 298–299]). Für $l, m, n \in \mathbb{N}$ mit $l \neq 0$ und $l \cdot m = l \cdot n$ ist $m = n$.

Beweis. Da $l \cdot m = l \cdot n$ ist, gilt $l \cdot m \leq l \cdot n$, was nach Proposition 1.5.5 bedeutet, dass $m \leq n$ ist. Ebenso gilt $l \cdot m \geq l \cdot n$, woraus $m \geq n$ folgt. Wegen der Antisymmetrie ist daher $m = n$. $\square$

Ebenfalls können wir nun zeigen, dass es bezüglich der Multiplikation im Allgemeinen keine Inversen gibt.

Proposition 1.5.7. Seien $n, n^{-1} \in \mathbb{N}$ mit $n \cdot n^{-1} = 1$. Dann ist $n = 1$.

Beweis. Die Fälle $n = 0$ und $n^{-1} = 0$ lassen sich wegen Lemma 1.4.9 ausschließen. Falls $n = 1$ ist, sind wir fertig.

Angenommen, es gäbe ein $n \geq 2$ und ein $n^{-1} \in \mathbb{N}$, sodass $n \cdot n^{-1} = 1$ ist. Dann ist auch $n^{-1} \geq 2$, da $n \cdot 1 = n \neq 1$ wäre. Es ist somit $n \cdot n^{-1} \geq 2 \cdot 2 = 4 > 1$ – ein Widerspruch zu $n \cdot n^{-1} = 1$. Somit bleibt nur $n = n^{-1} = 1$. $\square$

1.5.1 Die natürliche Ordnung als Wohlordnung

Im Verlauf dieser Diplomarbeit werden Schritt für Schritt die Zahlenbereiche erweitert. Bei jeder dieser Erweiterungen erhalten wir nützliche Eigenschaften, müssen für diese aber auch mit einer Eigenschaft des vorigen Zahlenbereichs „bezahlen“ [34, S. 30].

Die natürlichen Zahlen zeichnen sich gegenüber allen anderen Zahlenbereichen durch eine besondere Eigenschaft aus: Bei der natürlichen Ordnung handelt es sich um eine *Wohlordnung*. Das bedeutet, dass jede nichtleere Teilmenge der natürlichen Zahlen ein kleinstes Element bezüglich $\leq$ besitzt. Diese Eigenschaft geht bereits beim Übergang zu den ganzen Zahlen verloren, da beispielsweise die Menge der negativen ganzen Zahlen nichtleer, aber nach unten unbeschränkt ist.

Prinzipiell ist die Wohlordbarkeit keine Eigenschaft, die ausschließlich für die natürlichen Zahlen gilt. Nach dem Wohlordnungssatz lässt sich für *jede* Menge eine Wohlordnungsrelation finden (vgl. [19, S. 238]). Es ist allerdings nicht garantiert, dass diese

Ordnungsrelation mit der algebraischen Struktur der Menge verträglich ist. Somit ließe sich auch die Menge der ganzen Zahlen wohlordnen – Rechengesetze, wie sie im vorigen Abschnitt bewiesen wurden, müssen dann aber nicht gültig sein.

Für den Beweis, dass es sich bei $\leq$ um eine Wohlordnung handelt, benötigen wir eine spezielle Form der vollständigen Induktion: die *Ordnungsinduktion*.

Proposition 1.5.8 ([34, S. 3]). *Sei $\varphi(n)$ eine Aussage für natürliche Zahlen. Falls $\varphi(0)$ erfüllt ist, und für alle $n \in \mathbb{N}$ aus $\varphi(0) \wedge \varphi(1) \wedge \dots \wedge \varphi(n)$ die Gültigkeit von $\varphi(n+1)$ folgt, dann ist $\varphi(n)$ wahr für alle $n \in \mathbb{N}$.*

Beweis. Für $n \in \mathbb{N}$ sei $\Phi(n) := \bigwedge_{i=0}^n \varphi(i)$, sowie $N := \{n \in \mathbb{N} \mid \Phi(n)\}$. Dann ist $0 \in N$, da $\varphi(0)$ per Voraussetzung erfüllt ist.

Sei $n \in N$. Dann ist $\Phi(n) = \varphi(0) \wedge \varphi(1) \wedge \dots \wedge \varphi(n)$ wahr, weswegen nach Voraussetzung $\varphi(n+1)$ gilt. Damit ist $\Phi(n+1) = \Phi(n) \wedge \varphi(n+1)$ ebenfalls wahr, also $n+1 \in N$. Aufgrund des Induktionsprinzips folgt daher, dass $N = \mathbb{N}$ ist. $\square$

Satz 1.5.9 ([34, S. 7]). *Sei M eine nichtleere Teilmenge von $\mathbb{N}$. Dann besitzt M ein kleinstes Element: $\exists m \in M : \forall k \in M : m \leq k$.*

Beweis. Angenommen, M hätte kein kleinstes Element. Dann sei $S := \mathbb{N} \setminus M$ die Menge aller natürlichen Zahlen, die nicht in M enthalten sind. Es ist jedenfalls $0 \in S$, da sonst 0 nach Proposition 1.5.3 das kleinste Element von M wäre.

Sei $n \in S$. Angenommen, für alle $k \in \mathbb{N}$ mit $k \leq n$ wäre $k \in S$. Dann kann $n+1$ nicht in M sein, da es sonst ein Minimum der Menge wäre – denn alle natürlichen Zahlen bis einschließlich n sind nicht in M enthalten. Somit ist $n+1 \in S$. Aufgrund von Proposition 1.5.8 folgt, dass $S = \mathbb{N}$ ist. Da $S = \mathbb{N} \setminus M$ ist, muss aber $M = \emptyset$ sein – ein Widerspruch zu $M \neq \emptyset$.

Somit besitzt M ein kleinstes Element. $\square$

2 Die ganzen Zahlen

Betrachten wir nochmals Proposition 1.4.14. Deren Aussage war, dass in der Gleichung $n + m = 0$ über den natürlichen Zahlen beide Summanden gleich Null sein müssen. Daraus haben wir abgeleitet, dass Null die einzige natürliche Zahl ist, die über ein additives Inverses verfügt.

Aus algebraischer Sicht ist dieser Zustand nicht ideal. In vielen Fällen ist es notwendig, Gleichungen der Form $x + a = b$ zu lösen, was bei einer fehlenden Umkehroperation für die Addition nicht allgemein möglich ist. Wir suchen daher nach einer Erweiterung der natürlichen Zahlen um additive Inverse: die ganzen Zahlen. Die Frage nach der Invertierbarkeit der Multiplikation führt im nächsten Kapitel schließlich auf die rationalen Zahlen.

Die ganzen Zahlen stellen in dieser Arbeit die erste Zahlenbereichserweiterung dar. Eine wichtige Frage, die sich in diesem Kontext ergibt, ist, wie wir formal unsere bisherige Notation der Zahlen sowie Operatoren auf den neuen Zahlenbereich übertragen können. Eine Antwort darauf wird in Abschnitt 2.4 gegeben.

Wie im weiteren Verlauf dieses Kapitels gezeigt wird, handelt es sich bei den ganzen Zahlen zusammen mit der Addition und Multiplikation um einen nullteilerfreien Ring, jedoch um keinen Körper. Diese Struktur motiviert insbesondere Fragen zur Teilbarkeit, was in den ganzen Zahlen zu den Primzahlen und dem Hauptsatz der elementaren Zahlentheorie führt.

2.1 Geschichte der negativen Zahlen

Wir behandeln zwar in dieser Arbeit die ganzen Zahlen, und damit insbesondere die negativen Zahlen, vor den rationalen Zahlen, doch war diese Entwicklung historisch gesehen genau umgekehrt. Die Verwendung positiver rationaler Zahlen ist bereits im alten Ägypten belegt; davon zeugen etwa der Moskauer Papyrus (um 1850 v. Chr.) und der Papyrus Rhind (um 1650 v. Chr.) [56, S. 113–114, 116; 5, S. 39].

Negative Zahlen wurden vermutlich das erste Mal im zweiten Jahrhundert v. Chr. in China verwendet [40, S. 90]. Bereits damals war eine Form des Dezimalsystems in Verwendung, in dem Zahlen mit Hilfe von Stäbchen auf einem Raster dargestellt wurden, wobei jeder Spalte einer Zehnerpotenz entsprach [31, S. 148–152; 56, S. 52–53]. Um negative Zahlen von positiven Zahlen zu unterscheiden, wurden Stäbchen verschiedener Farben verwendet: Rot für positive Zahlen, und Schwarz für negative Zahlen [40, S. 90]. Regeln für das Rechnen mit negativen Zahlen werden bereits in der *Jiu Zhang Suanshu*, einem der ältesten chinesischen Werke über Mathematik, erwähnt [40, S. 90]. Die historische Datierung dieses Werks ist schwierig; vermutlich stammt es aus dem zweiten bis ersten Jahrhundert v. Chr. [29, S. 82].



Abbildung 2.1: Ein Rechenbrett aus einer japanischen Zeichnung [58]

Das Konzept der negativen Zahlen wurde im 7. Jahrhundert n. Chr. vom indischen Mathematiker Brahmagupta aufgegriffen, der in einem seiner Werke ebenfalls Regeln für das Rechnen mit diesen angibt [56, S. 102]. Im 12. Jahrhundert sah der indische Mathematiker Bhaskara erstmals negative Lösungen von Gleichungen als zulässig an,

während solche Gleichungen noch von Diophantos von Alexandria „sinnlos“ genannt wurden [29, S. 111; 40, S. 90].

Bis negative Zahlen in Europa Akzeptanz fanden, dauerte es viele Jahrhunderte. Es ist unklar, wann sie den europäischen Mathematikern bekannt wurden, allerdings wurden sie lange Zeit skeptisch bis ablehnend betrachtet. Leonardo von Pisa (auch als Fibonacci bekannt), war einer der ersten, die mit negativen Zahlen arbeiteten: In seinen *Liber Abaci* (Anfang des 13. Jahrhunderts) gibt er Beispiele mit negativen Lösungen an, die er im finanziellen Kontext als Schulden interpretiert hat. Es dauerte allerdings bis ins 16. Jahrhundert, bis negative Lösungen von Gleichungen von Girolamo Cardano als zulässig erkannt wurden [5, S. 321]; in dem Zusammenhang stieß Cardano auch erstmals auf die komplexen Zahlen (siehe Abschnitt 5.1).

Im 17. Jahrhundert „bewies“ der englische Mathematiker John Wallis, dass jede negative Zahl kleiner als Null und größer als Plus Unendlich ist. Sogar Leonhard Euler fühlte noch im 18. Jahrhundert die Notwendigkeit zu erklären, warum Minus mal Minus gleich Plus ist [39, S. 14].

2.2 Konstruktion der ganzen Zahlen

Wie in der Einführung dieses Kapitels angesprochen, liegt das Bestreben bei der Konstruktion der ganzen Zahlen darin, Gleichungen der Form $x + a = b$ immer lösen zu können. Das ist nur dann allgemein möglich, wenn es zu jeder Zahl a ein additives Inverses $-a$ gibt. Ziel ist also, die Struktur natürlichen Zahlen mit der Addition auf eine Gruppe zu erweitern.

Betrachten wir, naiv, die Zahl -2 . Wir können diese auf verschiedene Weisen darstellen: $-2 = 0 - 2 = 1 - 3 = 2 - 4 = \dots$. Jedes dieser Minuend-Subtrahend-Zahlenpaare stellt die Zahl -2 dar. Es würde sich also prinzipiell anbieten, die Paare $(0,2)$, $(1,3)$, $(2,4)$, ... mit -2 zu identifizieren, sie als äquivalent bezüglich einer (noch festzulegenden) Relation anzusehen.

Betrachten wir die Gleichung $0 - 2 = 1 - 3$. Durch Umformen erhalten wir $0 + 3 = 1 + 2$. Auch die Gleichung $1 - 3 = 2 - 4$ lässt sich durch Umformen auf die Gestalt $1 + 4 = 2 + 3$

bringen. Allgemein gesprochen: Haben wir die Zahlenpaare (a, b) sowie (c, d) , so wollen wir diese als äquivalent ansehen, falls $a - b = c - d$ ist. Das ist aber nach obigen Überlegungen äquivalent zu $a + d = c + b$. Da dieser Ausdruck nur natürliche Zahlen und die Addition enthält, können wir folgende Äquivalenzrelation auf $\mathbb{N} \times \mathbb{N}$ definieren.

Definition 2.2.1 ([34, S. 14]). Seien $(a, b), (c, d) \in \mathbb{N} \times \mathbb{N}$. Dann definieren wir die Relation $\sim$ auf $\mathbb{N} \times \mathbb{N}$ durch

$$(a, b) \sim (c, d) :\Leftrightarrow a + d = c + b.$$

Proposition 2.2.2. Bei der Relation $\sim$ auf $\mathbb{N} \times \mathbb{N}$ handelt es sich um eine Äquivalenzrelation. Es sind daher folgende Eigenschaften erfüllt:

- (1) $\forall a \in \mathbb{N} \times \mathbb{N} : a \sim a$ (Reflexivität).
- (2) $\forall a, b \in \mathbb{N} \times \mathbb{N} : a \sim b \Rightarrow b \sim a$ (Symmetrie).
- (3) $\forall a, b, c \in \mathbb{N} \times \mathbb{N} : (a \sim b \wedge b \sim c) \Rightarrow a \sim c$ (Transitivität).

Beweis. Seien $a = (a_1, a_2), b = (b_1, b_2)$ und $c = (c_1, c_2) \in \mathbb{N} \times \mathbb{N}$.

Zur Reflexivität: Wegen $a_1 + a_2 = a_1 + a_2$ ist $a \sim a$.

Zur Symmetrie: Angenommen, es wäre $a \sim b$, also $a_1 + b_2 = b_1 + a_2$. Lesen wir diese Gleichung von rechts nach links, so erhalten wir direkt $b \sim a$.

Zur Transitivität: Angenommen, $a \sim b$ und $b \sim c$, also $a_1 + b_2 = b_1 + a_2$ und $b_1 + c_2 = c_1 + b_2$. Addieren der Gleichungen liefert

$$(a_1 + b_2) + (b_1 + c_2) = (b_1 + a_2) + (c_1 + b_2),$$

woraus mit Hilfe der Kürzungsregel und der Kommutativität der Addition $a_1 + c_2 = c_1 + a_2$ folgt. Es ist damit $a \sim c$. □

Die obige Äquivalenzrelation fasst, wie zuvor gezeigt, die Paare $(0,2), (1,3), (2,4), \dots$ zu einer Äquivalenzklasse zusammen, die wir als $C(0,2)$ anschreiben. Da die Äquivalenzklasse mehrere Repräsentanten besitzt, ist $C(0,2) = C(1,3) = C(2,4) = \dots$. Diese

Mehrdeutigkeit wird in den nachfolgenden Beweisen oft die Frage nach der Wohldefiniertheit von Aussagen über, oder Operationen auf den Äquivalenzklassen aufwerfen.

Die ganzen Zahlen definieren wir nun als die Faktormenge von $\mathbb{N} \times \mathbb{N}$ bezüglich der Relation $\sim$.

Definition 2.2.3 ([34, S. 14]). Die Menge der ganzen Zahlen $\mathbb{Z}$ ist definiert als

$$\mathbb{Z} := \mathbb{N} \times \mathbb{N} / \sim = \{C(a, b) \mid a, b \in \mathbb{N}\}.$$

Abbildung 2.2 stellt die Äquivalenzklassen in den ganzen Zahlen grafisch dar. Dabei werden die Paare aus $\mathbb{N} \times \mathbb{N}$ als Punkte in einem kartesischen Koordinatensystem verstanden. Die Äquivalenzklassen werden durch die Geraden (beziehungsweise Strahlen) dargestellt, welche parallel zueinander verlaufen. Dort, wo sie die x -Achse schneiden, befinden sich die natürlichen, beziehungsweise die neu dazugekommenen negativen ganzen Zahlen [34, S. 15; 48, S. 152].

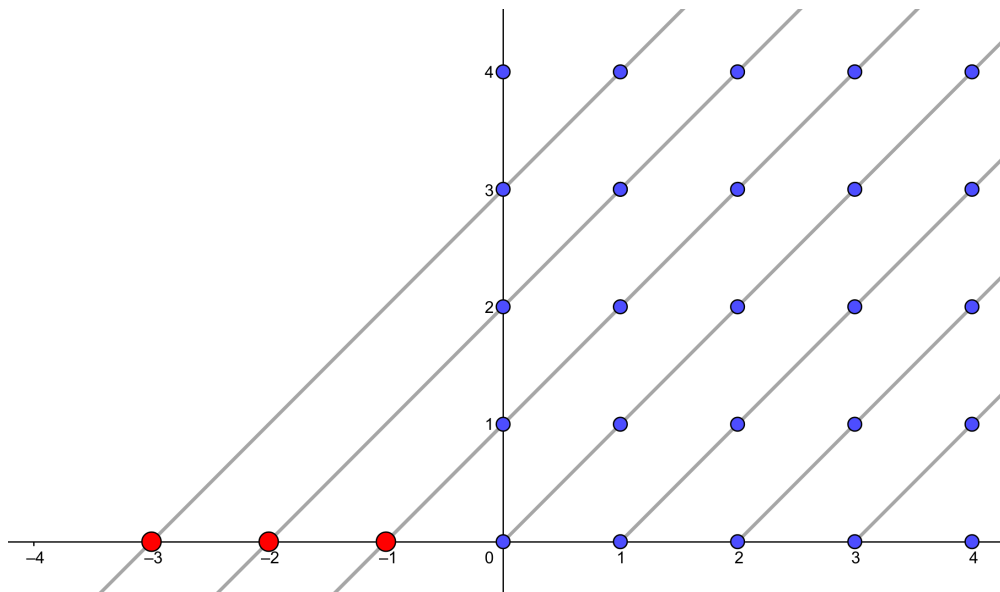


Abbildung 2.2: Die blauen Punkte sind jeweils die Zahlenpaare aus $\mathbb{N} \times \mathbb{N}$, während die Geraden/Strahlen die Elemente der Äquivalenzklassen verbinden.

Folgende Proposition präzisiert diese grafische Beschreibung der Äquivalenzklassen.

Proposition 2.2.4. Sei $C(a, b) \in \mathbb{Z}$ und $n \in \mathbb{N}$. Dann ist $C(a, b) = C(a + n, b + n)$.

Beweis. Wir wollen zeigen, dass $(a, b) \sim (a + n, b + n)$ erfüllt ist. Durch Anwenden von Definition 2.2.1 erhalten wir $a + (b + n) = (a + n) + b$, was offensichtlich stimmt. Da zwei Äquivalenzklassen stets disjunkt oder gleich sind, und $(a + n, b + n) \in C(a, b)$ ist, folgt die Behauptung. $\square$

In Abbildung 2.2 erkennt man, dass sich jeder ganzen Zahl genau eine natürliche Zahl zuordnen lässt, welche sich auf der x -Achse, beziehungsweise auf der y -Achse befindet. Das wird vor allem für die Überlegungen zur Einbettung der natürlichen Zahlen in die ganzen Zahlen wichtig sein.

Proposition 2.2.5. *Sei $C(a, b) \in \mathbb{Z}$. Dann gibt es genau ein $n \in \mathbb{N}$, sodass $C(a, b) = C(n, 0)$ oder $C(a, b) = C(0, n)$ ist.*

Beweis. Ist $a \geq b$, so gibt es nach Definition 1.5.1 ein $n \in \mathbb{N}$ mit $a = b + n$; dieses ist wegen der Kürzungsregel eindeutig. Es ist dann $(n, 0) \sim (a, b)$, da $n + b = a + 0$, und somit $C(a, b) = C(n, 0)$.

Ist hingegen $a \leq b$, dann gibt es ein eindeutiges $n \in \mathbb{N}$ mit $b = a + n$. Es ist $(0, n) \sim (a, b)$, da $0 + b = a + n$, woraus $C(a, b) = C(0, n)$ folgt. $\square$

Wir zeigen außerdem folgendes Lemma, das sich direkt aus Abbildung 2.2 ergibt. Es sagt aus, dass die Elemente von $C(0, 0)$ genau jene Punkte im Gitter sind, die sich auf der ersten Hauptdiagonale befinden.

Lemma 2.2.6. *Sei $C(a, b) \in \mathbb{Z}$. Dann ist $C(a, b) = C(0, 0)$ genau dann, wenn $a = b$ ist.*

Beweis. Angenommen, $C(a, b) = C(0, 0)$, dann ist $(a, b) \sim (0, 0)$, was nach Definition 2.2.1 bedeutet, dass $a + 0 = 0 + b$ ist, also $a = b$. Ist umgekehrt $a = b$, dann ist auch $a + 0 = 0 + b$, und daher $(a, b) \sim (0, 0)$. $\square$

2.3 Rechenoperationen auf $\mathbb{Z}$

Im folgenden Abschnitt wollen wir in den ganzen Zahlen die bekannte Addition und Multiplikation definieren. Vorab überlegen wir: Die Äquivalenzklasse $C(a, b)$ stellt die ganze Zahl $a - b$ dar (etwa $C(0, 2)$ die Zahl -2). Wollen wir nun zwei ganze Zahlen $a - b$ und $c - d$ addieren, erhalten wir $(a + c) - (b + d)$, was der Äquivalenzklasse $C(a + c, b + d)$ entspricht. Die Addition in den ganzen Zahlen wird somit, ähnlich wie in einem K^n -Vektorraum über einem Körper K , komponentenweise durchgeführt.

Definition 2.3.1. Die Addition $+$: $\mathbb{Z} \times \mathbb{Z} \rightarrow \mathbb{Z}$ ist definiert durch

$$C(a, b) + C(c, d) := C(a + c, b + d).$$

Die Definition der Addition wurde über Repräsentanten von Restklassen durchgeführt. Wir müssen daher zeigen, dass die Addition wohldefiniert, also unabhängig von der Wahl der Repräsentanten ist.

Proposition 2.3.2. Die Addition auf den ganzen Zahlen ist wohldefiniert. Sind also (a, b) und (a', b') Repräsentanten der Restklasse $C(a, b)$, sowie (c, d) und (c', d') Repräsentanten von $C(c, d)$, dann ist $C(a, b) + C(c, d) = C(a', b') + C(c', d')$.

Beweis. Nach Definition 2.2.1 ist $a + b' = a' + b$ und $c + d' = c' + d$. Durch Addieren dieser Gleichungen und Umformen erhalten wir

$$(a + c) + (b' + d') = (a' + c') + (b + d),$$

was nach Definition 2.2.1 bedeutet, dass $(a + c, b + d) \sim (a' + c', b' + d')$ ist. Es ist daher $(a + c, b + d) \in C(a' + c', b' + d')$. Da zwei Äquivalenzklassen bezüglich einer Äquivalenzrelation stets entweder disjunkt oder gleich sind, folgt daraus $C(a + c, b + d) = C(a' + c', b' + d')$, was äquivalent zu $C(a, b) + C(c, d) = C(a', b') + C(c', d')$ ist. $\square$

Wie bereits bei der Einführung in dieses Kapitel erwähnt, war es die Absicht bei der Erweiterung auf die ganzen Zahlen, Gleichungen der Form $x + a = b$ stets lösen zu können. Das ist möglich, weil es sich bei $(\mathbb{Z}, +)$ um eine Gruppe handelt.

Proposition 2.3.3. *Bei $(\mathbb{Z}, +, C(0,0))$ handelt es sich um eine kommutative Gruppe. Es sind daher folgende Eigenschaften erfüllt:*

- (1) $\forall a, b, c \in \mathbb{Z} : (a + b) + c = a + (b + c)$ (Assoziativität).
- (2) $\forall a, b \in \mathbb{Z} : a + b = b + a$ (Kommutativität).
- (3) $\forall a \in \mathbb{Z} : a + C(0,0) = a$ (Nullelement).
- (4) $\forall a \in \mathbb{Z} : \exists -a \in \mathbb{Z} : a + (-a) = C(0,0)$ (Invertierbarkeit).

Beweis. Seien $a = (a_1, a_2)$, $b = (b_1, b_2)$ und $c = (c_1, c_2)$ ganze Zahlen. Für den Beweis der Assoziativität verwenden wir die entsprechende Eigenschaft in den natürlichen Zahlen.

$$\begin{aligned}
 (a + b) + c &= (C(a_1, a_2) + C(b_1, b_2)) + C(c_1, c_2) = \\
 &= C(a_1 + b_1, a_2 + b_2) + C(c_1, c_2) = \\
 &= C((a_1 + b_1) + c_1, (a_2 + b_2) + c_2) = \\
 &= C(a_1 + (b_1 + c_1), a_2 + (b_2 + c_2)) = \\
 &= C(a_1, a_2) + C(b_1 + c_1, b_2 + c_2) = \\
 &= C(a_1, a_2) + (C(b_1, b_2) + C(c_1, c_2)).
 \end{aligned}$$

Es ist zu beachten, dass in den vorigen Schritten mit dem Operator $+$ zwei verschiedene Additionen bezeichnet wurden, nämlich einerseits die Addition in den ganzen Zahlen (Definition 2.3.1), und andererseits die Addition in den natürlichen Zahlen (Definition 1.4.1). Da aus dem Kontext jedoch klar ist, um welche der beiden Additionen es sich handelt, dürfen wir uns diese Ungenauigkeit erlauben.

Die Kommutativität können wir ebenfalls auf die natürlichen Zahlen zurückführen:

$$\begin{aligned}
 a + b &= C(a_1, a_2) + C(b_1, b_2) = \\
 &= C(a_1 + b_1, a_2 + b_2) = \\
 &= C(b_1 + a_1, b_2 + a_2) = \\
 &= C(b_1, b_2) + C(a_1, a_2) = \\
 &= b + a.
 \end{aligned}$$

Analog folgt, dass es sich bei $C(0, 0)$ um das Nullelement handelt:

$$a + C(0,0) = C(a_1, a_2) + C(0,0) = C(a_1 + 0, a_2 + 0) = C(a_1, a_2) = a.$$

Zuletzt beweisen wir die Existenz von Inversen. Wir setzen $-a := C(a_2, a_1)$. Dann ist

$$a + (-a) = C(a_1, a_2) + C(a_2, a_1) = C(a_1 + a_2, a_2 + a_1) = C(0, 0),$$

wobei wir für die letzte Gleichheit Lemma 2.2.6 verwendet haben. $\square$

Wir vereinbaren nun, für $k, m \in \mathbb{Z}$ anstelle von $k + (-m)$ kürzer $k - m$, und das additive Inverse von k als $-k$ zu schreiben.

Wie in den natürlichen Zahlen gilt auch in den ganzen Zahlen die Kürzungsregel für die Addition. Da wir nun jedoch eine Gruppenstruktur haben, lässt sich die Regel sehr einfach beweisen.

Proposition 2.3.4. *Seien $k, m, n \in \mathbb{Z}$. Falls $m + k = n + k$ gilt, dann ist $m = n$.*

Beweis. Die Aussage folgt direkt durch Addieren von $-k$ auf beiden Seiten der Gleichung. $\square$

Proposition 2.3.5. *Seien $a, b \in \mathbb{Z}$. Dann gibt es für die Gleichung $x + a = b$ stets eine eindeutige Lösung $x \in \mathbb{Z}$.*

Beweis. Setzen wir $x := b - a$, erhalten wir $x + a = (b - a) + a = b + (-a + a) = b + 0 = b$. Somit ist x eine Lösung der Gleichung.

Angenommen, x und x' erfüllen die Gleichung. Dann ist $x + a = b = x' + a$, woraus wegen der Kürzungsregel für die Addition $x = x'$ folgt. $\square$

Um die Multiplikation in den ganzen Zahlen zu definieren, überlegen wir wieder: Sind $a - b$ und $c - d$ ganze Zahlen, dann ist $(a - b) \cdot (c - d) = (a \cdot c + b \cdot d) - (a \cdot d + b \cdot c)$. Nach den Überlegungen, die zu Definition 2.2.1 geführt haben, entspricht das der ganzen Zahl $C(a \cdot c + b \cdot d, a \cdot d + b \cdot c)$. Wir definieren daher die Multiplikation wie folgt.

Definition 2.3.6. Die Multiplikation $\cdot : \mathbb{Z} \times \mathbb{Z} \rightarrow \mathbb{Z}$ ist definiert durch

$$C(a, b) \cdot C(c, d) := C(a \cdot c + b \cdot d, a \cdot d + b \cdot c).$$

Wegen der Verwendung von Repräsentanten von Restklassen in der Definition müssen wir zeigen, dass die Multiplikation wohldefiniert ist.

Proposition 2.3.7. *Die Multiplikation auf den ganzen Zahlen ist wohldefiniert: Sind $(a, b), (a', b') \in C(a, b)$ sowie $(c, d), (c', d') \in C(c, d)$, so ist $C(a, b) \cdot C(c, d) = C(a', b') \cdot C(c', d')$.*

Beweis. Nach Definition 2.2.1 ist $(a, b) \sim (a', b')$ und $(c, d) \sim (c', d')$, also $a + b' = a' + b$ und $c + d' = c' + d$.

Multiplizieren wir die erste Gleichung jeweils einmal mit c und d , so erhalten wir

$$\begin{aligned} a \cdot c + b' \cdot c &= a' \cdot c + b \cdot c, \\ a \cdot d + b' \cdot d &= a' \cdot d + b \cdot d. \end{aligned}$$

Vertauschen wir nun die Seiten der zweiten Gleichung und addieren beide Gleichungen, so erhalten wir

$$a \cdot c + b' \cdot c + a' \cdot d + b \cdot d = a' \cdot c + b \cdot c + a \cdot d + b' \cdot d. \quad (2.1)$$

Nun multiplizieren wir die zweite Gleichung jeweils einmal mit a' und b' , und erhalten

$$\begin{aligned} c \cdot a' + d' \cdot a' &= c' \cdot a' + d \cdot a', \\ c \cdot b' + d' \cdot b' &= c' \cdot b' + d \cdot b'. \end{aligned}$$

Wir vertauschen wieder die Seiten der zweiten Gleichung und addieren beide Gleichungen. Dadurch erhalten wir

$$c \cdot a' + d' \cdot a' + c' \cdot b' + d \cdot b' = c' \cdot a' + d \cdot a' + c \cdot b' + d' \cdot b'. \quad (2.2)$$

Betrachten wir nun die Gleichungen 2.1 und 2.2, umgeordnet und Teile durch Klammern hervorgehoben.

$$\begin{aligned} a \cdot c + b \cdot d + (a' \cdot d + b' \cdot c) &= (a' \cdot c + b' \cdot d) + a \cdot d + b \cdot c \\ a' \cdot d' + b' \cdot c' + (a' \cdot c + b' \cdot d) &= a' \cdot c' + b' \cdot d' + (a' \cdot d + b' \cdot c) \end{aligned}$$

Die durch Klammern hervorgehobenen Terme haben ein Gegenstück auf der gegenüberliegenden Seite in der jeweils anderen Gleichung. Addieren wir nun beide Gleichungen, so können wir die Kürzungsregel verwenden, um diese Terme zu streichen. Dadurch erhalten wir folgende Gleichung.

$$(a \cdot c + b \cdot d) + (a' \cdot d' + b' \cdot c') = (a' \cdot c' + b' \cdot d') + (a \cdot d + b \cdot c)$$

Nach Definition 2.2.1 ist daher

$$(a \cdot c + b \cdot d, a \cdot d + b \cdot c) \sim (a' \cdot c' + b' \cdot d', a' \cdot d' + b' \cdot c').$$

Somit erhalten wir

$$\begin{aligned} C(a, b) \cdot C(c, d) &= C(a \cdot c + b \cdot d, a \cdot d + b \cdot c) = \\ &= C(a' \cdot c' + b' \cdot d', a' \cdot d' + b' \cdot c') = \\ &= C(a', b') \cdot C(c', d'). \end{aligned} \quad \square$$

Die ganzen Zahlen bilden die erste Ringstruktur in unserer Konstruktion der Zahlenbereiche.

Proposition 2.3.8. $(\mathbb{Z}, +, \cdot, C(0,0), C(1,0))$ ist ein kommutativer Ring mit Eins. Es sind daher folgende Eigenschaften erfüllt:

- (1) $(\mathbb{Z}, +, C(0,0))$ ist eine kommutative Gruppe.
- (2) $\forall a, b, c \in \mathbb{Z} : (a \cdot b) \cdot c = a \cdot (b \cdot c)$ (Assoziativität).
- (3) $\forall a, b \in \mathbb{Z} : a \cdot b = b \cdot a$ (Kommutativität).
- (4) $\forall a \in \mathbb{Z} : a \cdot C(1,0) = a$ (Einselement).

(5) $\forall a, b, c \in \mathbb{Z} : a \cdot (b + c) = a \cdot b + a \cdot c$ (Distributivität).

Beweis. Seien $a = (a_1, a_2)$, $b = (b_1, b_2)$ und $c = (c_1, c_2)$ ganze Zahlen. Die erste Behauptung wurde bereits in Proposition 2.3.3 gezeigt.

Zur besseren Übersicht schreiben wir die Komponenten der Äquivalenzklassen untereinander an. Die Assoziativität lässt sich dann beweisen, indem wir die Assoziativität und Distributivität in den natürlichen Zahlen verwenden.

$$\begin{aligned}
 (a \cdot b) \cdot c &= \left(C \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot C \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} \right) \cdot C \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} = \\
 &= C \begin{pmatrix} a_1 \cdot b_1 + a_2 \cdot b_2 \\ a_1 \cdot b_2 + a_2 \cdot b_1 \end{pmatrix} \cdot C \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} = \\
 &= C \begin{pmatrix} (a_1 \cdot b_1 + a_2 \cdot b_2) \cdot c_1 + (a_1 \cdot b_2 + a_2 \cdot b_1) \cdot c_2 \\ (a_1 \cdot b_1 + a_2 \cdot b_2) \cdot c_2 + (a_1 \cdot b_2 + a_2 \cdot b_1) \cdot c_1 \end{pmatrix} = \\
 &= C \begin{pmatrix} a_1 \cdot (b_1 \cdot c_1 + b_2 \cdot c_2) + a_2 \cdot (b_1 \cdot c_2 + b_2 \cdot c_1) \\ a_1 \cdot (b_1 \cdot c_2 + b_2 \cdot c_1) + a_2 \cdot (b_1 \cdot c_1 + b_2 \cdot c_2) \end{pmatrix} = \\
 &= C \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot C \begin{pmatrix} b_1 \cdot c_1 + b_2 \cdot c_2 \\ b_1 \cdot c_2 + b_2 \cdot c_1 \end{pmatrix} = \\
 &= C \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot \left(C \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} \cdot C \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} \right) = \\
 &= a \cdot (b \cdot c).
 \end{aligned}$$

Die Kommutativität lässt sich ebenfalls unter Verwendung entsprechender Eigenschaften der natürlichen Zahlen beweisen.

$$\begin{aligned}
 a \cdot b &= (a_1, a_2) \cdot (b_1, b_2) = \\
 &= (a_1 \cdot b_1 + a_2 \cdot b_2, a_1 \cdot b_2 + a_2 \cdot b_1) = \\
 &= (b_1 \cdot a_1 + b_2 \cdot a_2, b_1 \cdot a_2 + b_2 \cdot a_1) = \\
 &= (b_1, b_2) \cdot (a_1, a_2) = \\
 &= b \cdot a.
 \end{aligned}$$

Zum Einselement: Es ist

$$a \cdot (1, 0) = (a_1, a_2) \cdot (1, 0) = (a_1 \cdot 1 + a_2 \cdot 0, a_1 \cdot 0 + a_2 \cdot 1) = (a_1, a_2) = a.$$

Zuletzt beweisen wir die Gültigkeit des Distributivgesetzes.

$$\begin{aligned} a \cdot (b + c) &= C(a_1, a_2) \cdot (C(b_1, b_2) + C(c_1, c_2)) = \\ &= C(a_1, a_2) \cdot C(b_1 + c_1, b_2 + c_2) = \\ &= C(a_1 \cdot (b_1 + c_1) + a_2 \cdot (b_2 + c_2), a_1 \cdot (b_2 + c_2) + a_2 \cdot (b_1 + c_1)) = \\ &= C(a_1 \cdot b_1 + a_1 \cdot c_1 + a_2 \cdot b_2 + a_2 \cdot c_2, a_1 \cdot b_2 + a_1 \cdot c_2 + a_2 \cdot b_1 + a_2 \cdot c_1) = \\ &= C(a_1 \cdot b_1 + a_2 \cdot b_2, a_1 \cdot b_2 + a_2 \cdot b_1) + C(a_1 \cdot c_1 + a_2 \cdot c_2, a_1 \cdot c_2 + a_2 \cdot c_1) = \\ &= C(a_1, a_2) \cdot C(b_1, b_2) + C(a_1, a_2) \cdot C(c_1, c_2) = \\ &= a \cdot b + a \cdot c. \end{aligned} \quad \square$$

Proposition 2.3.9. *Die ganzen Zahlen sind nullteilerfrei.*

Beweis. Angenommen, für $m, n \in \mathbb{Z}$ ist $m \cdot n = C(0, 0)$. Nach Proposition 2.2.5 gibt es $\hat{m}, \hat{n} \in \mathbb{N}$ mit $m = C(\hat{m}, 0)$ oder $m = C(0, \hat{m})$, und $n = C(\hat{n}, 0)$ oder $n = C(0, \hat{n})$. Das Produkt $m \cdot n$ kann damit folgende Formen annehmen:

- $C(\hat{m}, 0) \cdot C(\hat{n}, 0) = C(\hat{m} \cdot \hat{n}, 0),$
- $C(\hat{m}, 0) \cdot C(0, \hat{n}) = C(0, \hat{m} \cdot \hat{n}),$
- $C(0, \hat{m}) \cdot C(\hat{n}, 0) = C(0, \hat{m} \cdot \hat{n}),$
- $C(0, \hat{m}) \cdot C(0, \hat{n}) = C(\hat{m} \cdot \hat{n}, 0).$

Da $m \cdot n = C(0, 0)$ ist, folgt wegen Lemma 2.2.6, dass $\hat{m} \cdot \hat{n} = 0$ ist. Wegen Proposition 1.4.13 erhalten wir $\hat{m} = 0$ oder $\hat{n} = 0$, also $m = C(0, 0)$ oder $n = C(0, 0)$. $\square$

Wie in den natürlichen Zahlen gilt auch in den ganzen Zahlen die Kürzungsregel für die Multiplikation.

Proposition 2.3.10. *Seien $k, m, n \in \mathbb{Z}$ mit $k \neq C(0, 0)$. Dann folgt aus $m \cdot k = n \cdot k$, dass $m = n$ ist.*

Beweis. Angenommen, es ist $m \cdot k = n \cdot k$, dann können wir diese Gleichung zu $(m - n) \cdot k = C(0,0)$ umformen. Wegen Proposition 2.3.9 ist $m - n = C(0,0)$, also $m = n$. $\square$

Um eine Ordnung auf den ganzen Zahlen zu definieren, gehen wir ähnlich vor wie in den natürlichen Zahlen. Wir „wissen“, dass eine ganze Zahl k kleiner gleich einer ganzen Zahl m ist, falls k am Zahlenstrahl links von m liegt. In dem Fall können wir m erreichen, indem wir von k aus $m - k$ Schritte nach rechts gehen – was nichts anderes bedeutet, als dass wir zu k die positive Zahl $m - k$ addieren.

Damit wir diese Überlegung für die Definition verwenden können, benötigen wir zunächst den Begriff der positiven und negativen ganzen Zahlen. Dieser wird durch Abbildung 2.2 und Proposition 2.2.5 motiviert.

Definition 2.3.11. Sei $m \in \mathbb{Z}$ mit $m \neq C(0,0)$ und $n \in \mathbb{N}$ die nach Proposition 2.2.5 existierende zugehörige natürliche Zahl. Wir nennen m positiv, falls $m = C(n,0)$ ist. Ist hingegen $m = C(0, n)$, so nennen wir die Zahl negativ. $C(0,0)$ wird als weder positiv noch negativ definiert.

Zudem definieren wir

$$\begin{aligned}\mathbb{Z}^+ &:= \{C(k,0) \mid k \in \mathbb{N} \setminus \{0\}\}, \\ \mathbb{Z}^- &:= \{C(0, k) \mid k \in \mathbb{N} \setminus \{0\}\}\end{aligned}$$

als die Menge der positiven, beziehungsweise der negativen ganzen Zahlen.

Proposition 2.3.12. *Es ist $\mathbb{Z} = \mathbb{Z}^- \cup \{C(0,0)\} \cup \mathbb{Z}^+$ und $\mathbb{Z}^- \cap \mathbb{Z}^+ = \emptyset$. Jede ganze Zahl ist daher entweder positiv, negativ, oder gleich dem Nullelement.*

Beweis. Offensichtlich ist $C(0,0)$ in der Vereinigung enthalten. Sei $m \in \mathbb{Z} \setminus \{C(0,0)\}$ und n die nach Proposition 2.2.5 existierende eindeutige zugehörige natürliche Zahl. Es gilt dann entweder $m = C(n,0)$ oder $m = C(0, n)$, was per Definition $m \in \mathbb{Z}^+$, beziehungsweise $m \in \mathbb{Z}^-$ bedeutet. Somit ist auch m Element der Vereinigung.

Gäbe es eine ganze Zahl, die sowohl positiv wie auch negativ wäre, so wäre $C(a,0) = C(0, b)$, was nach Definition 2.2.1 bedeutet, dass $a + b = 0$ ist. Nach Proposition 1.4.14

ist damit $a = b = 0$, und die gesuchte ganze Zahl daher gleich dem Nullelement – ein Widerspruch dazu, dass die Zahl sowohl positiv als auch negativ ist. Somit kann eine ganze Zahl nur entweder positiv, negativ, oder gleich dem Nullelement sein. $\square$

Aus obiger Proposition folgt insbesondere $\mathbb{Z} \setminus \mathbb{Z}^- = \mathbb{Z}^+ \cup \{C(0,0)\}$ und $\mathbb{Z} \setminus \mathbb{Z}^+ = \mathbb{Z}^- \cup \{C(0,0)\}$.

Betrachten wir $k, m \in \mathbb{Z}$. Sind sowohl k als auch m positiv (ist also $k = C(\hat{k}, 0)$ und $m = C(\hat{m}, 0)$) dann ist auch $k + m = C(\hat{k} + \hat{m}, 0)$ positiv. Analog können wir zeigen, dass die Summe zweier negativer Zahlen wieder negativ ist.

Stellen wir ähnliche Überlegungen für die Multiplikation an, erhalten wir die aus der Schule bekannten Rechenregeln „Plus mal Minus ist Minus“ und „Minus mal Minus ist Plus“. Denn ist k positiv und m negativ, dann ist $k \cdot m = C(\hat{k}, 0) \cdot C(0, \hat{m}) = C(0, \hat{k} \cdot \hat{m})$ eine negative ganze Zahl. Sind hingegen sowohl k als auch m negativ, so ist $k \cdot m = C(0, \hat{k}) \cdot C(0, \hat{m}) = C(\hat{k} \cdot \hat{m}, 0)$ positiv.

Lemma 2.3.13. *Sei $k \in \mathbb{Z}$. Falls k positiv ist, dann ist $-k$ negativ. Ist umgekehrt k negativ, dann ist $-k$ positiv.*

Beweis. Angenommen, k ist positiv, also $k = C(n, 0)$ für ein $n \in \mathbb{N} \setminus \{0\}$. Dann ist $-k = C(0, n)$, da $C(n, 0) + C(0, n) = C(n, n) = C(0, 0)$ ist. Nach Definition 2.3.11 ist $-k$ negativ.

Falls k negativ ist, also $k = C(0, n)$ für ein $n \in \mathbb{N} \setminus \{0\}$, dann ist wegen der selben Überlegung $-k = C(n, 0)$ eine positive ganze Zahl. $\square$

Definition 2.3.14. Seien $k, m \in \mathbb{Z}$. Dann definieren wir

$$k \leq m :\Leftrightarrow \exists n \in \mathbb{Z} \setminus \mathbb{Z}^- : k + n = m.$$

Wir schreiben zudem $k \geq m$ für $m \leq k$. Ist das obige n positiv (also ungleich $C(0,0)$), können wir $k < m$ beziehungsweise $m > k$ schreiben.

Wenig überraschend ist, dass die negativen ganzen Zahlen kleiner als $C(0,0)$, und die positiven ganzen Zahlen größer als $C(0,0)$ sind. Denn für eine negative ganze Zahl

$C(0, n)$ ist $C(n, 0)$ positiv und $C(0, n) + C(n, 0) = C(n, n) = C(0, 0)$, und für eine positive ganze Zahl $C(n, 0)$ offensichtlich $C(0, 0) + C(n, 0) = C(n, 0)$.

Proposition 2.3.15. *Bei der Relation $\leq$ auf $\mathbb{Z}$ handelt es sich um eine Totalordnung. Es gilt daher $k \leq m$ oder $m \leq k$ für alle $k, m \in \mathbb{Z}$. Zusätzlich besitzt $\leq$ die Eigenschaften einer Ordnungsrelation:*

- (1) $\forall k \in \mathbb{Z} : k \leq k$ (Reflexivität).
- (2) $\forall k, m \in \mathbb{Z} : (k \leq m \wedge m \leq k) \Rightarrow k = m$ (Antisymmetrie).
- (3) $\forall k, m, n \in \mathbb{Z} : (k \leq m \wedge m \leq n) \Rightarrow k \leq n$ (Transitivität).

Beweis. Zur Reflexivität: Für $k \in \mathbb{Z}$ ist $k + C(0, 0) = k$, und daher $k \leq k$.

Zur Antisymmetrie: Falls $k \leq m$ und $m \leq k$, dann gibt es $n, n' \in \mathbb{Z} \setminus \mathbb{Z}^-$ mit $k + n = m$ und $m + n' = k$. Einsetzen der ersten Gleichung in die zweite liefert $k + (n + n') = k$, also $n + n' = C(0, 0)$. Wären n oder n' positiv, so wäre die Summe $n + n'$ ebenfalls positiv – im Widerspruch dazu, dass $n + n' = C(0, 0)$ ist. Es kann somit nur $n = n' = C(0, 0)$ sein, woraus $k = m$ folgt.

Zur Transitivität: Ist $k \leq m$ und $m \leq n$, so gibt es nichtnegative $a, b \in \mathbb{Z}$ mit $k + a = m$ und $m + b = n$. Einsetzen der ersten Gleichung in die zweite liefert $k + (a + b) = n$, also $k \leq n$.

Wir wollen nun zeigen, dass stets zwei ganze Zahlen über diese Relation verglichen werden können. Seien $k, m \in \mathbb{Z}$. Die Gleichung $k + x = m$ hat wegen Proposition 2.3.5 garantiert eine eindeutige Lösung $x \in \mathbb{Z}$. Zudem ist $m - x = k$. Falls $x = C(0, 0)$ ist, dann ist $k = m$, und daher auch $k \leq m$. Andernfalls ist nach Lemma 2.3.13 x oder $-x$ positiv, woraus entweder $k \leq m$ oder $m \leq k$ folgt. $\square$

Wir können jetzt insbesondere folgern, dass alle negativen Zahlen kleiner als alle positiven Zahlen sind. Denn nach unserer vorigen Überlegung ist für $k \in \mathbb{Z}^-$ stets $k \leq C(0, 0)$, und für $m \in \mathbb{Z}^+$ immer $C(0, 0) \leq m$ erfüllt. Aufgrund der Transitivität ist daher $k \leq m$.

Proposition 2.3.16. *Die Relation $\leq$ auf $\mathbb{Z}$ erfüllt die Ordnungsaxiome:*

- (1) Für $k, l, m \in \mathbb{Z}$ ist $k \leq m \Leftrightarrow k + l \leq m + l$.

(2) Für $k, m \in \mathbb{Z}$ mit $k > C(0,0)$ und $m > C(0,0)$ ist $k \cdot m > C(0,0)$.

Beweis. Zu (1): Angenommen $k \leq m$, dann gibt es nach Definition 2.3.14 ein nichtnegatives $n \in \mathbb{Z}$ mit $k+n = m$. Wegen der Kürzungsregel ist das äquivalent zu $(k+l)+n = m+l$, was $k+l \leq m+l$ bedeutet.

Zu (2): Sind k und m positiv, so ist $k = C(\hat{k},0)$ und $m = C(\hat{m},0)$ mit $\hat{k} \neq 0$ und $\hat{m} \neq 0$. Es ist dann $k \cdot m = C(\hat{k},0) \cdot C(\hat{m},0) = C(\hat{k} \cdot \hat{m},0)$, also auch $k \cdot m$ positiv, da die natürlichen Zahlen nullteilerfrei sind. $\square$

Mit Ungleichungen in den ganzen Zahlen dürfen wir fast wie gewohnt rechnen. Die einzige Ausnahme besteht beim Multiplizieren mit oder Kürzen von negativen Faktoren in einer Ungleichung. In diesem Fall dreht sich das Ungleichheitszeichen um.

Proposition 2.3.17. Seien $k, l, m, n \in \mathbb{Z}$.

- (1) Falls $k \leq l$ und $m \leq n$, dann ist $k + m \leq l + n$.
- (2) Ist k positiv, dann ist $l \leq m \Leftrightarrow l \cdot k \leq m \cdot k$.
- (3) Ist k negativ, dann ist $l \leq m \Leftrightarrow l \cdot k \geq m \cdot k$.

Beweis. Zu (1): Ist $k \leq l$ und $m \leq n$, so gibt es nichtnegative $a, b \in \mathbb{Z}$ mit $k + a = l$ und $m + b = n$. Addieren dieser Gleichungen liefert $(k + m) + (a + b) = l + n$, was $k + m \leq l + n$ bedeutet.

Zu (2): Angenommen, $l \leq m$, dann gibt es ein nichtnegatives $a \in \mathbb{Z}$ mit $l + a = m$. Multiplizieren der Gleichung mit k und Anwenden des Distributivgesetzes liefert $l \cdot k + a \cdot k = m \cdot k$. Da a und k nichtnegativ sind, ist nach Proposition 2.3.16 auch $a \cdot k$ nichtnegativ (die Multiplikation mit Null ergibt, wie in jedem Ring, Null). Daher ist $l \cdot k \leq m \cdot k$.

Sei umgekehrt $l \cdot k \leq m \cdot k$. Dann ist $l \cdot k + (m - l) \cdot k = m \cdot k$ und es gibt ein $a \in \mathbb{Z} \setminus \mathbb{Z}^-$ mit $l \cdot k + a = m \cdot k$. Wegen der Kürzungsregel erhalten wir, dass $a = (m - l) \cdot k$ eine nichtnegative Zahl ist. Da k positiv ist, muss $m - l$ nichtnegativ sein, also $l \leq m$.

Zu (3): Angenommen, $l \leq m$ und $k \in \mathbb{Z}$ ist negativ, dann ist wegen Lemma 2.3.13 $-k$ positiv. Nach (2) ist daher $l \leq m$ äquivalent zu $l \cdot (-k) \leq m \cdot (-k)$. Das ist wiederum

äquivalent dazu, dass es ein nichtnegatives $a \in \mathbb{Z}$ gibt, sodass $l \cdot (-k) + a = m \cdot (-k)$ ist. Multiplizieren wir beide Seiten mit $-C(1,0)$, so erhalten wir $l \cdot k - a = m \cdot k$ (diese Rechenregeln gelten in jedem Ring mit Eins). Addieren wir nun auf beiden Seiten der Gleichung a , so erhalten wir $l \cdot k = m \cdot k + a$, was äquivalent ist zu $m \cdot k \leq l \cdot k$. $\square$

Auch in den ganzen Zahlen gibt es im Allgemeinen keine multiplikativen Inversen.

Proposition 2.3.18. *Seien $k, k^{-1} \in \mathbb{Z}$ mit $k \cdot k^{-1} = C(1,0)$. Dann ist $k = C(1,0)$ oder $k = C(0,1)$.*

Beweis. Falls $k \cdot k^{-1} = C(1,0)$ ist, so müssen k und k^{-1} beide entweder positiv oder negativ sein, da das Produkt eines positiven und negativen Faktors stets negativ ist. Es gibt also $\hat{k}, \hat{k}^{-1} \in \mathbb{N}$ mit $k = C(\hat{k}, 0)$ und $k^{-1} = C(\hat{k}^{-1}, 0)$, oder $k = C(0, \hat{k})$ und $k^{-1} = C(0, \hat{k}^{-1})$. In beiden Fällen ist $C(1,0) = k \cdot k^{-1} = C(\hat{k} \cdot \hat{k}^{-1}, 0)$, also $1 = \hat{k} \cdot \hat{k}^{-1}$. Aus Proposition 1.5.7 folgt, dass $\hat{k} = \hat{k}^{-1} = 1$ ist. Somit ist $k = k^{-1} = C(1,0)$ oder $k = k^{-1} = C(0,1)$. $\square$

Wir können nun auch die Betragsfunktion definieren, welche die bekannten Eigenschaften besitzt (Dreiecksungleichung, Produktregel).

Definition 2.3.19. Sei $k \in \mathbb{Z}$. Dann definieren wir

$$|k| := \begin{cases} k & \text{wenn } k \geq C(0), \\ -k & \text{sonst.} \end{cases}$$

In den natürlichen Zahlen war die natürliche Ordnung zugleich eine Wohlordnung (siehe Satz 1.5.9). Das ist bei der Ordnung auf den ganzen Zahlen nicht mehr der Fall. Betrachten wir beispielsweise die Menge der negativen ganzen Zahlen $\mathbb{Z}^-$, so besitzt diese kein kleinstes Element. Denn angenommen $C(0, m) \in \mathbb{Z}^-$ wäre minimal, dann wäre $C(0, m+1) \in \mathbb{Z}^-$ und $C(0, m+1) \leq C(0, m)$, da $C(0, m+1) + C(1,0) = C(1, m+1) = C(0, m)$, wobei im letzten Schritt Proposition 2.2.4 verwendet wurde – im Widerspruch zur Minimalität von $C(0, m)$. Somit gibt es in $\mathbb{Z}^-$ kein kleinstes Element, weswegen $\leq$ keine Wohlordnung sein kann.

2.4 Einbettung von $\mathbb{N}$ in $\mathbb{Z}$

Aus der Schulmathematik „wissen“ wir, dass $\mathbb{N} \subseteq \mathbb{Z}$ ist. Betrachten wir jedoch Definition 1.3.3 und Definition 2.2.3, so sehen wir, dass das offensichtlich nicht der Fall sein kann. Dennoch verwenden wir im Alltag eine Notation, in der wir die nichtnegativen ganzen Zahlen mit den natürlichen Zahlen identifizieren, und für die negativen ganzen Zahlen den positiven ganzen Zahlen ein Minus voranstellen. In diesem Abschnitt wollen wir diese übliche Alltagsnotation mathematisch fundieren.

Die Grundidee, welche hinter der Einbettung der natürlichen in die ganzen Zahlen steckt, ist mit dem vorigen Absatz bereits angesprochen: Wir wollen die nichtnegativen ganzen Zahlen mit den natürlichen Zahlen identifizieren. Für jedes $n \in \mathbb{N}$ suchen wir genau ein $m \in \mathbb{Z} \setminus \mathbb{Z}^-$, sodass n „möglichst gut“ zu m passt. „Möglichst gut“ bedeutet, dass Eigenschaften, die wir für $\mathbb{N}$ bewiesen haben, auch für $\mathbb{Z} \setminus \mathbb{Z}^-$ gelten. Ist also beispielsweise in den natürlichen Zahlen $2 \cdot 2 = 4$, soll das auch für die zu 2 und 4 zugehörigen ganzen Zahlen gelten.

Algebraisch formuliert bedeutet das: Die Zuordnung zwischen $\mathbb{N}$ und $\mathbb{Z}$ soll strukturverträglich sein. Insbesondere sollen die Strukturen um $\mathbb{N}$ und $\mathbb{Z} \setminus \mathbb{Z}^-$ zueinander isomorph sein. Ist die Isomorphie erfüllt, so können wir aus algebraischer Sicht tatsächlich davon sprechen, dass es sich nur um eine Umbenennung der Elemente handelt [50, S. 230]. Wir dürfen also ungenau sein, und die Notationen vermischen.

Definition 2.4.1 ([vgl. 50, S. 262; 34, S. 13, 16]). Seien $(R_1, +, \cdot)$ und $(R_2, \oplus, \otimes)$ zwei Halbringe. Wir nennen eine injektive Abbildung $\varphi : R_1 \rightarrow R_2$ eine Einbettung, falls folgende Eigenschaften erfüllt sind:

- (1) $\forall a, b \in R_1 : \varphi(a + b) = \varphi(a) \oplus \varphi(b).$
- (2) $\forall a, b \in R_1 : \varphi(a \cdot b) = \varphi(a) \otimes \varphi(b).$

Da es sich bei dieser Abbildung um einen Homomorphismus handelt, gelten für sie auch dieselben Eigenschaften (vgl. [50, S. 233] und Lemma 3.4.1).

Um nun die natürlichen Zahlen in die ganzen Zahlen einzubetten, verwenden wir die Idee, die uns bereits zur Definition der positiven ganzen Zahlen geführt hat: Wir bilden

die natürliche Zahl n auf die ganze Zahl $C(n,0)$ ab und zeigen, dass damit die vorigen Eigenschaften erfüllt sind [34, S. 16].

Proposition 2.4.2. *Sei $\varphi : \mathbb{N} \rightarrow \mathbb{Z}$ mit $\varphi(n) := C(n, 0)$. Dann handelt es sich bei φ um eine Einbettung.*

Beweis. Seien $n, m \in \mathbb{N}$. Dann ist $\varphi(n + m) = C(n + m, 0) = C(n, 0) + C(m, 0) = \varphi(n) + \varphi(m)$. Somit ist die erste Eigenschaft erfüllt.

Zudem ist $\varphi(n \cdot m) = C(n \cdot m, 0) = C(n, 0) \cdot C(m, 0) = \varphi(n) \cdot \varphi(m)$.

Um die Injektivität zu beweisen, nehmen wir an, es wäre $\varphi(n) = \varphi(m)$. Es ist dann $C(n, 0) = C(m, 0)$, woraus wir wegen Definition 2.4.1 direkt $n = m$ erhalten. $\square$

Weiters gilt, dass sich die Ordnung der natürlichen Zahlen auf die ganzen Zahlen überträgt.

Proposition 2.4.3. *Sei φ die Einbettung aus Proposition 2.4.2. Dann ist für $m, n \in \mathbb{N}$ $n \leq m \Leftrightarrow \varphi(n) \leq \varphi(m)$.*

Beweis. Nehmen wir $n \leq m$ an. Dann gibt es ein $k \in \mathbb{N}$ mit $n + k = m$. Es ist daher $\varphi(n) + \varphi(k) = \varphi(m)$, was $\varphi(n) \leq \varphi(m)$ bedeutet, da $\varphi(k)$ nichtnegativ ist.

Ist umgekehrt $\varphi(n) \leq \varphi(m)$, so gibt es ein $C(k, 0) \in \mathbb{Z}$, sodass $\varphi(n) + C(k, 0) = \varphi(m)$ ist. Wegen $C(k, 0) = \varphi(k)$ ist $m = n + k$, also $n \leq m$. $\square$

Wir haben somit erfolgreich gezeigt, dass die Struktur der natürlichen Zahlen mit jener der ganzen Zahlen verträglich ist. Insbesondere sind die natürlichen Zahlen isomorph zu den nichtnegativen ganzen Zahlen.

Um die im Alltag übliche Notation zu erhalten, vereinbaren wir, künftig jede nichtnegative ganze Zahl $C(n, 0)$ als n anzuschreiben, etwa $C(2, 0)$ als 2. Da wir festgestellt haben, dass $C(0, n) = -C(n, 0)$ gilt, schreiben wir eine negative ganze Zahl $C(0, n)$ in Zukunft als $-n$ an. Für das Nullelement $C(0, 0)$ vereinbaren wir die übliche Notation ohne Vorzeichen [34, S. 16].

Somit ist $\mathbb{Z} = \{\dots, -3, -2, -1, 0, 1, 2, 3, \dots\}$. Obwohl es mathematisch gesehen ungenau ist, verwenden wir die Notationen $\mathbb{N} \subseteq \mathbb{Z}$ und $\mathbb{N} = \mathbb{Z} \setminus \mathbb{Z}^-$ um auszudrücken, dass die natürlichen Zahlen in den ganzen Zahlen eingebettet sind.

2.5 Der Hauptsatz der elementaren Zahlentheorie

Die ganzen Zahlen bilden nach Proposition 2.3.18 keinen Körper, da es von Null verschiedene Zahlen gibt, die kein multiplikatives Inverses besitzen. Das bedeutet insbesondere, dass die Gleichung $a \cdot x = b$ über den ganzen Zahlen im Allgemeinen nicht lösbar ist. Dennoch gibt es Gleichungen dieser Gestalt, die eine Lösung besitzen, etwa $2 \cdot x = 8$ mit $x = 4$. In der Schule schreiben wir dazu auch $8 : 2 = 4$, gesprochen als „Acht durch Zwei ergibt Vier“. Es sind Fragen rund um die Lösbarkeit solche Gleichungen, der Teilbarkeit, mit denen sich die elementare Zahlentheorie beschäftigt.

Im folgenden Abschnitt entwickeln wir die elementare Zahlentheorie so weit, dass wir einen wichtigen Satz über die ganzen Zahlen, den Hauptsatz der elementaren Zahlentheorie, beweisen können.

Definition 2.5.1. Seien $k, m \in \mathbb{Z}$ mit $k \neq 0$. Wir schreiben $k|m$, gesprochen als „ k teilt m “, falls es ein $n \in \mathbb{Z}$ mit $m = k \cdot n$ gibt.

Bemerkung. Falls zwischen zwei ganzen Zahlen Teilbarkeit gegeben ist, so gilt wegen der Kürzungsregel für die Multiplikation, dass n eindeutig ist.

In der Definition wurde Null als Teiler anderer Zahlen ausgeschlossen, denn für $m \in \mathbb{Z} \setminus \{0\}$ gibt es kein $l \in \mathbb{Z}$, sodass $0 \cdot l = m$ ist. Für $m = 0$ erfüllen jedoch alle ganzen Zahlen diese Eigenschaft – warum schließt man also $0|0$ aus? Der Grund dafür liegt darin, dass das n aus der Definition eindeutig sein soll, was nur dann der Fall ist, wenn man Null nicht zulässt (vgl. Proposition 2.5.3)

Die Teilbarkeitsrelation ist reflexiv, da $m \cdot 1 = m$ für $m \in \mathbb{Z}$ gilt. Sie ist auch transitiv, denn falls $k|m$ und $m|n$, so gibt es nach Definition $a, b \in \mathbb{Z}$, sodass $m = k \cdot a$ und $n = m \cdot b$ ist. Einsetzen der ersten Gleichung in die zweite liefert $n = k \cdot (a \cdot b)$, was $k|n$ bedeutet. Sie ist jedoch im Allgemeinen weder symmetrisch (es ist etwa $2|4$, aber nicht $4|2$) noch antisymmetrisch ($2|-2$ und $-2|2$, aber $2 \neq -2$).

Die Teilbarkeitsrelation induziert eine Form von „Ordnung“ auf den ganzen Zahlen ohne Null, wie folgende Proposition zeigt.

Proposition 2.5.2. *Seien $k, m \in \mathbb{Z} \setminus \{0\}$ mit $k|m$. Dann ist $|k| \leq |m|$.*

Beweis. Da k ein Teiler von m ist, gibt es ein $n \in \mathbb{Z}$, sodass $m = k \cdot n$ ist. Da $m \neq 0$ ist, muss wegen der Nullteilerfreiheit der ganzen Zahlen auch $n \neq 0$ sein. Wenden wir auf beiden Seiten der Gleichung die Betragsfunktion sowie deren Rechenregeln an, erhalten wir $|m| = |k| \cdot |n|$.

Es ist $1 \leq |n|$, woraus mit Proposition 2.3.17 folgt, dass $|k| \leq |k| \cdot |n| = |m|$ ist, also $|k| \leq |m|$. $\square$

Aus $m|a$ und $m|b$ folgt $m|(a+b)$, denn ist $a = m \cdot k$ und $b = m \cdot l$, so ist $a+b = m \cdot (k+l)$. Ebenso folgt daraus $m|(a \cdot b)$, denn es ist $a \cdot b = m \cdot (k \cdot m \cdot l)$.

Ein wichtiges Resultat der Zahlentheorie ist die Möglichkeit der Division mit Rest.

Proposition 2.5.3. *Seien $k, m \in \mathbb{Z}$ mit $k \neq 0$. Dann gibt es eindeutige $q, r \in \mathbb{Z}$ mit $0 \leq r < |k|$, sodass $m = k \cdot q + r$ ist.*

Beweis. Wir zeigen zunächst, dass der Satz für alle natürlichen m erfüllt ist und folgern anschließend daraus, dass er daher auch für alle negativen ganzen Zahlen gelten muss.

Um zu beweisen, dass für alle natürlichen Zahlen eine Division mit Rest durch k möglich ist, verwenden wir das Induktionsprinzip. Für $m = 0$ ist die Aussage offensichtlich wahr. Sei daher m eine natürliche Zahl, die sich in der Form $m = k \cdot q + r$ anschreiben lässt, wobei $0 \leq r < |k|$ ist. Addieren wir auf beiden Seiten der Gleichung Eins, so erhalten wir

$$m + 1 = k \cdot q + (r + 1).$$

Ist $r + 1 < |k|$, so sind wir fertig. Andernfalls ist $r + 1 = |k|$, und wir können wegen der Distributivität k herausheben, wobei wir $|k| = k \cdot \operatorname{sgn} k$ verwenden:

$$m + 1 = k \cdot (q + \operatorname{sgn} k) + 0.$$

Wir können daher auch für $m + 1$ eine Division mit Rest durch k durchführen, und wegen des Induktionsprinzips somit für alle natürlichen Zahlen.

Ist m eine negative ganze Zahl, so ist $-m$ wegen Lemma 2.3.13 positiv. Aufgrund des vorigen Schrittes gibt es $q, r \in \mathbb{Z}$ mit $0 \leq r < |k|$, sodass $-m = k \cdot q + r$ ist. Multiplizieren wir diese Gleichung mit -1 , erhalten wir $m = k \cdot (-q) - r$. Addieren und Subtrahieren von $|k|$ auf der rechten Seite der Gleichung liefert wegen $|k| = k \cdot \operatorname{sgn} k$

$$m = k \cdot (-q - \operatorname{sgn}(k)) + (|k| - r).$$

Da $0 \leq |k| - r < |k|$ ist, haben wir somit die Existenz von q und r für die negativen ganzen Zahlen gezeigt.

Zum Beweis der Eindeutigkeit: Angenommen, es gäbe $q_1, q_2, r_1, r_2 \in \mathbb{Z}$ mit $0 \leq r_1 < |k|$ und $0 \leq r_2 < |k|$, sodass $m = k \cdot q_1 + r_1$ und $m = k \cdot q_2 + r_2$ ist. Durch Subtrahieren und Umformen der beiden Gleichungen erhalten wir $r_2 - r_1 = k \cdot (q_1 - q_2)$. Angenommen es wäre $q_1 \neq q_2$, dann gilt $k | (r_2 - r_1)$, woraus nach Proposition 2.5.2 folgt, dass $|k| \leq |r_1 - r_2|$ ist. Es ist aber $0 \leq r_1 < |k|$ und $0 \leq r_2 < |k|$, also $|r_1 - r_2| < |k|$ – ein Widerspruch. Somit bleibt nur, dass $q_1 = q_2$ ist, woraus $r_1 = r_2$ folgt. $\square$

Bereits im antiken Griechenland wurde untersucht, welche Teiler verschiedene Zahlen besitzen. So können wir beispielsweise die Teilmengen T_m ganzer Zahlen m bestimmen:

$$\begin{aligned} T_1 &= \{\pm 1\} \\ T_2 &= \{\pm 1, \pm 2\} \\ T_3 &= \{\pm 1, \pm 3\} \\ T_4 &= \{\pm 1, \pm 2, \pm 4\} \\ T_5 &= \{\pm 1, \pm 5\} \\ T_6 &= \{\pm 1, \pm 2, \pm 3, \pm 6\} \\ &\vdots \end{aligned}$$

Anhand dieser Auflistung lässt sich leicht erkennen, dass die Teilermenge einer Zahl

immer ± 1 und die Zahl selbst enthält. Die Schnittmenge der Teilmengen zweier Zahlen, sprich die Menge ihrer gemeinsamen Teiler, ist somit stets nichtleer. Interessant ist in dem Zusammenhang das Maximum solcher Schnittmengen, das wir den *größten gemeinsamen Teiler* nennen. Dieser existiert aufgrund dieser Beobachtungen und Proposition 2.5.2 jedenfalls, und ist größer gleich Eins.

Definition 2.5.4. Seien $m_1, m_2, \dots, m_n \in \mathbb{Z} \setminus \{0\}$. Wir definieren

$$\text{ggT}(m_1, m_2, \dots, m_n) := \max \bigcap_{i=1}^n T_n.$$

Während manche Zahlen mehrere Teiler besitzen, gibt es auch Zahlen, die ausschließlich durch ± 1 und $\pm$ sich selbst teilbar sind. Diese Zahlen sind als Primzahlen bekannt.

Definition 2.5.5. Eine Zahl $n \in \mathbb{N} \setminus \{1\}$ heißt Primzahl, wenn sie ausschließlich durch ± 1 und $\pm n$ teilbar ist: $k|n \Rightarrow (|k| = 1 \vee |k| = n)$.

Warum Eins als Primzahl ausgeschlossen wird, hat zum Teil historische Hintergründe. So wurde etwa lange Zeit „Eins“ nicht als Zahl angesehen; Euklid definierte „Zahl“ als Vielheit von Einheiten. Würde man Eins als Primzahl betrachten, müsste man einige Sätze der Zahlentheorie, wie etwa den Hauptsatz, dahingehend umformulieren, dass mit „Primzahl“ eine „Primzahl ungleich Eins“ gemeint ist. Der einfachere Weg ist, Eins bereits im Voraus als Primzahl auszuschließen [6].

Primzahlen bilden gewissermaßen die „Atome“, aus denen die natürlichen Zahlen zusammengesetzt sind. Der Hauptsatz der Zahlentheorie besagt, dass sich jede natürliche Zahl mit Ausnahme von Null und Eins eindeutig als Produkt von Primzahlen darstellen lässt. Beispiele für solche Darstellungen wären etwa $16 = 2 \cdot 2 \cdot 2 \cdot 2$ oder $105 = 3 \cdot 5 \cdot 7$.

Um den Hauptsatz der Zahlentheorie beweisen zu können, benötigen wir zunächst folgende Propositionen.

Proposition 2.5.6 ([35, S. 80–81]). Seien $m, n \in \mathbb{Z} \setminus \{0\}$ mit $\text{ggT}(m, n) = 1$. Dann gibt es $a, b \in \mathbb{Z}$, sodass $a \cdot m + b \cdot n = 1$ ist.

Beweis. Wir definieren zunächst $M := \{a \cdot m + b \cdot n \mid a, b \in \mathbb{Z}\}$. Dann ist $|m| + |n| \in M$ und $|m| + |n| > 0$, da beide Zahlen ungleich Null sind. Es ist somit $M \cap \mathbb{Z}^+$ eine nichtleere Teilmenge der natürlichen Zahlen. Aus Satz 1.5.9 folgt, dass es ein minimales Element $e \in M \cap \mathbb{Z}^+$ gibt, welches sich nach Definition von M als $e = a_0 \cdot m + b_0 \cdot n$ mit $a_0, b_0 \in \mathbb{Z}$ darstellen lässt.

Nach Proposition 2.5.3 gibt es $k, r \in \mathbb{Z}$ mit $0 \leq r < e$, sodass $m = k \cdot e + r$ ist. Setzen wir den vorigen Ausdruck für e in diese Gleichung ein, so erhalten wir $m = k \cdot (a_0 \cdot m + b_0 \cdot n) + r$, was sich zu $r = (1 - k \cdot a_0) \cdot m + (-k \cdot b_0) \cdot n$ umformen lässt. Es ist also $r \in M$.

Angenommen, $r \neq 0$, dann wäre $r < e$ ein Widerspruch zur Minimalität von e . Daher muss $r = 0$ sein, und e ein Teiler von m . Analog können wir zeigen, dass e ein Teiler von n ist. Daher ist e ein gemeinsamer Teiler von m und n , und da $e \in \mathbb{N}$ sowie $\text{ggT}(m, n) = 1$ ist, muss $e = 1$ sein. Es ist somit $a_0 \cdot m + b_0 \cdot n = 1$. $\square$

Das folgende Lemma geht auf Euklid zurück, der es in seinen *Elementen* bewies. Wir beweisen hier zunächst eine allgemeinere Version dieses Lemmas, das den größten gemeinsamen Teiler verwendet. Die Originalaussage Euklids über Primzahlen folgt als Korollar.

Lemma 2.5.7. Seien $k, m, n \in \mathbb{Z}$ mit $k \mid (m \cdot n)$ und $\text{ggT}(k, m) = 1$. Dann ist $k \mid n$.

Beweis. Nach Proposition 2.5.6 gibt es $a_0, b_0 \in \mathbb{Z}$, sodass $a_0 \cdot k + b_0 \cdot m = 1$ ist. Multiplizieren dieser Gleichung mit n liefert $a_0 \cdot k \cdot n + b_0 \cdot m \cdot n = n$. Da $k \mid a_0 \cdot k \cdot n$ und $k \mid b_0 \cdot m \cdot n$ gilt, muss k auch die Summe teilen, also $k \mid n$. $\square$

Korollar 2.5.7.1. Seien $m, n \in \mathbb{Z}$ und p eine Primzahl, sodass $p \mid (m \cdot n)$. Dann gilt $p \mid m$ oder $p \mid n$.

Beweis. Falls $p \mid m$, so sind wir fertig. Andernfalls ist $\text{ggT}(p, m) = 1$, da es sich bei p um eine Primzahl handelt. Aus Lemma 2.5.7 folgt damit $p \mid n$. $\square$

Satz 2.5.8. Sei $m \in \mathbb{N} \setminus \{0, 1\}$. Dann besitzt m eine bis auf Permutation eindeutige Zerlegung in Primfaktoren: $m = p_1 \cdot p_2 \cdot \dots \cdot p_n$ für nicht notwendigerweise verschiedene Primzahlen $p_1, \dots, p_n$.

Beweis. Zur Existenz: Angenommen, es gäbe Zahlen in $\mathbb{N} \setminus \{0,1\}$, die keine Zerlegung in Primfaktoren besitzen. Sei M die nichtleere Menge dieser Zahlen. Da $M \subseteq \mathbb{N}$ ist, gibt es nach Satz 1.5.9 ein minimales $m \in M$. Da m keine Primzahl sein kann (denn dann wäre die Zahl selbst ihre Primfaktorzerlegung), muss es $k, l \in \mathbb{N}$ mit $m = k \cdot l$ geben. Nach Proposition 2.5.2 ist $k < m$ und $l < m$. Da m die kleinste natürliche Zahl war, die keine Primfaktorzerlegung besitzt, müssen k und l eine solche Zerlegung besitzen. Deren Produkt ist dann eine Primfaktorzerlegung von m – ein Widerspruch dazu, dass m keine Primfaktorzerlegung besitzt. Somit lässt sich jede natürliche Zahl außer Null und Eins in Primfaktoren zerlegen.

Zur Eindeutigkeit: Angenommen,

$$m = p_1 \cdot p_2 \cdot \dots \cdot p_n = q_1 \cdot q_2 \cdot \dots \cdot q_k$$

für nicht notwendigerweise verschiedene Primzahlen p_i, q_j mit $1 \leq i \leq n$ und $1 \leq j \leq k$. Wir nehmen ohne Beschränkung der Allgemeinheit an, dass $k \leq n$ ist. Für $1 \leq j \leq k$ gilt $q_j | m = p_1 \cdot p_2 \cdot \dots \cdot p_n$. Nach Korollar 2.5.7.1 gibt es daher ein p_{a_j} , sodass $q_j | p_{a_j}$. Da es sich allerdings um Primzahlen handelt, muss $q_j = p_{a_j}$ gelten. Somit ist $q_1 = p_{a_1}$, $q_2 = p_{a_2}$, ..., $q_k = p_{a_k}$.

Wäre $k < n$, so könnte man in der Gleichung $p_1 \cdot p_2 \cdot \dots \cdot p_n = q_1 \cdot q_2 \cdot \dots \cdot q_k$ alle q_j ($1 \leq j \leq k$) kürzen, womit auf der rechten Seite der Gleichung Eins stünde, und auf der linken Seite ein nichtleeres Produkt von Zahlen größer als Eins – ein Widerspruch, genau wie der Fall $k > n$. Es ist somit $k = n$ und die Darstellung in Primfaktoren bis auf Permutation eindeutig. $\square$

3 Die rationalen Zahlen

Mit der Erweiterung auf die ganzen Zahlen haben wir erstmals einen nullteilerfreien Ring mit Eins erhalten. Während wir daher aktuell Gleichungen der Form $x + a = b$ stets lösen können, fehlt uns nach Proposition 2.3.18 im Allgemeinen diese Möglichkeit bei Gleichungen der Gestalt $a \cdot x + b = c$. Die Frage nach der allgemeinen Lösbarkeit solcher Gleichungen führt auf die rationalen Zahlen.

Die rationalen Zahlen sind sowohl algebraisch als auch historisch gesehen interessant. Einerseits bilden die rationalen Zahlen zum ersten Mal in der Kette der Zahlenbereichserweiterungen eine Körperstruktur, und sind daher aus algebraischem Blickwinkel in vielerlei Hinsicht ideal. Historisch gesehen sind die rationalen Zahlen deshalb spannend, weil sie bereits viele tausend Jahre vor den negativen ganzen Zahlen in Verwendung waren, etwa im alten Ägypten oder im antiken Griechenland, und deren Existenz weit weniger Kopferbrechen verursacht hat, als die negativen ganzen Zahlen (siehe Abschnitt 2.1). Diese historische Entwicklung wird auch im Schulunterricht nachgegangen: Während Bruchzahlen bereits in der vierten Schulstufe im Lehrplan vorgesehen sind, werden die negativen rationalen Zahlen erst in der siebten Schulstufe eingeführt [42, S. 157–158; 44, S. 6–7].

3.1 Konstruktion der rationalen Zahlen

Ziel der Erweiterung auf die rationalen Zahlen ist, Gleichungen der Form $a \cdot x + b = c$ für $a \neq 0$ stets lösen zu können. Formen wir diese Gleichung um, erhalten wir $a \cdot x = c - b$. Diese Gleichung ist dann stets lösbar, wenn es zu a ein multiplikatives Inverses a^{-1} gibt. Wir suchen also eine Erweiterung der ganzen Zahlen, in der es zu jedem Element $a \neq 0$ ein Inverses a^{-1} gibt, sodass $a \cdot a^{-1} = 1$ ist.

Die Erweiterung von den ganzen Zahlen auf die rationalen Zahlen geschieht ähnlich wie im letzten Kapitel. Wir überlegen zunächst anschaulich „naiv“: Da ein Bruch aus jeweils einer ganzen Zahl (dem Zähler) und einer positiven natürlichen Zahl (dem Nenner) besteht, bietet es sich an, dass wir die Bruchzahlen als Elemente aus $\mathbb{Z} \times (\mathbb{N} \setminus \{0\})$ verstehen. Nachdem die Bruchdarstellung nicht eindeutig ist (so ist etwa $\frac{1}{2} = \frac{2}{4} = \frac{3}{6} = \dots$), müssen wir verschiedene Elemente von $\mathbb{Z} \times (\mathbb{N} \setminus \{0\})$ miteinander über eine Äquivalenzrelation identifizieren. Die rationalen Zahlen definieren wir schließlich als die Faktormenge von $\mathbb{Z} \times (\mathbb{N} \setminus \{0\})$ über dieser Äquivalenzrelation.

Nehmen wir „naiv“ an, für zwei rationale Zahlen $\frac{a}{b}$ und $\frac{c}{d}$ wäre $\frac{a}{b} = \frac{c}{d}$. Multiplizieren wir mit b und d , so erhalten wir $a \cdot d = c \cdot b$. Da diese Gleichung ausschließlich natürliche und ganze Zahlen enthält, können wir sie zur Definition der gesuchten Äquivalenzrelation verwenden.

Definition 3.1.1 ([34, S. 18]). Seien $(a, b), (c, d) \in \mathbb{Z} \times (\mathbb{N} \setminus \{0\})$. Wir definieren die Relation $\sim$ auf dieser Menge als

$$(a, b) \sim (c, d) : \Leftrightarrow a \cdot d = c \cdot b.$$

Proposition 3.1.2. Die Relation $\sim$ auf $\mathbb{Z} \times (\mathbb{N} \setminus \{0\})$ ist eine Äquivalenzrelation. Es sind daher folgende Eigenschaften erfüllt:

- (1) $\forall a \in \mathbb{Z} \times (\mathbb{N} \setminus \{0\}) : a \sim a$ (Reflexivität).
- (2) $\forall a, b \in \mathbb{Z} \times (\mathbb{N} \setminus \{0\}) : a \sim b \Rightarrow b \sim a$ (Symmetrie).
- (3) $\forall a, b, c \in \mathbb{Z} \times (\mathbb{N} \setminus \{0\}) : (a \sim b \wedge b \sim c) \Rightarrow a \sim c$ (Transitivität).

Beweis. Seien $a = (a_1, a_2), b = (b_1, b_2)$ und $c = (c_1, c_2) \in \mathbb{Z} \times (\mathbb{N} \setminus \{0\})$. Die Reflexivität folgt direkt durch Einsetzen in die Definition: Es ist $a_1 \cdot a_2 = a_1 \cdot a_2$, also $a \sim a$.

Zur Symmetrie: Angenommen, es wäre $a \sim b$, also $a_1 \cdot b_2 = b_1 \cdot a_2$. Vertauschen wir die rechte Seite der Gleichung mit der linken, so folgt direkt $b \sim a$.

Zur Transitivität: Angenommen, es wäre $a \sim b$ und $b \sim c$. Dann ist $a_1 \cdot b_2 = b_1 \cdot a_2$ und $b_1 \cdot c_2 = c_1 \cdot b_2$. Multiplizieren wir die erste Gleichung mit c_2 und die zweite Gleichung

mit a_2 , so erhalten wir

$$a_1 \cdot b_2 \cdot c_2 = b_1 \cdot a_2 \cdot c_2$$

$$b_1 \cdot c_2 \cdot a_2 = c_1 \cdot b_2 \cdot a_2.$$

Wegen der Kommutativität der Multiplikation in den ganzen Zahlen ist also

$$a_1 \cdot b_2 \cdot c_2 = c_1 \cdot b_2 \cdot a_2,$$

woraus mit Hilfe der Kürzungsregel $a_1 \cdot c_2 = c_1 \cdot a_2$ folgt, also $a \sim c$. $\square$

Die Äquivalenzklasse $C(a, b)$ bezüglich dieser Relation schreiben wir auch als $\frac{a}{b}$ an, und nennen sie einen Bruch. Die erste Komponente a nennen wir den Nenner, und die zweite Komponente b den Zähler des Bruchs. Sind für zwei Brüche die Nenner gleich, so heißen diese gleichnamig.

Die Definition der rationalen Zahlen gelingt nun über die Faktormenge.

Definition 3.1.3 ([34, S. 18]). Die Menge der rationalen Zahlen $\mathbb{Q}$ ist definiert als

$$\mathbb{Q} := \mathbb{Z} \times (\mathbb{N} \setminus \{0\}) / \sim = \left\{ \frac{a}{b} \mid (a, b) \in \mathbb{Z} \times (\mathbb{N} \setminus \{0\}) \right\}.$$

Auch für die rationalen Zahlen können wir eine ähnliche grafische Darstellung wie für die ganzen Zahlen finden. In Abbildung 3.1 verbinden Geraden den Ursprung mit Punkten aus $\mathbb{Z} \times (\mathbb{N} \setminus \{0\})$, wobei Punkte derselben Äquivalenzklasse auf einer Geraden liegen. Der neue Zahlenstrahl ist die Gerade $f(x) = 1$, wobei die Schnittpunkte mit den Geraden der Äquivalenzklassen die rationalen Zahlen darstellen [34, S. 18].

Aus der Schule ist uns bekannt, dass wir Brüche erweitern und kürzen können, es gilt $\frac{a}{b} = \frac{a \cdot n}{b \cdot n}$ für alle $n \in \mathbb{N} \setminus \{0\}$. Die folgende Proposition beweist diese Umformungsregel.

Proposition 3.1.4. Sei $\frac{a}{b} \in \mathbb{Q}$ und $n \in \mathbb{N} \setminus \{0\}$. Dann ist $\frac{a \cdot n}{b \cdot n} = \frac{a}{b}$.

Beweis. Da es sich um Äquivalenzklassen handelt, müssen wir zeigen, dass $(a, b) \sim (a \cdot n, b \cdot n)$ ist. Das folgt allerdings direkt aus Definition 3.1.1, denn es ist $a \cdot b \cdot n = a \cdot n \cdot b$. $\square$

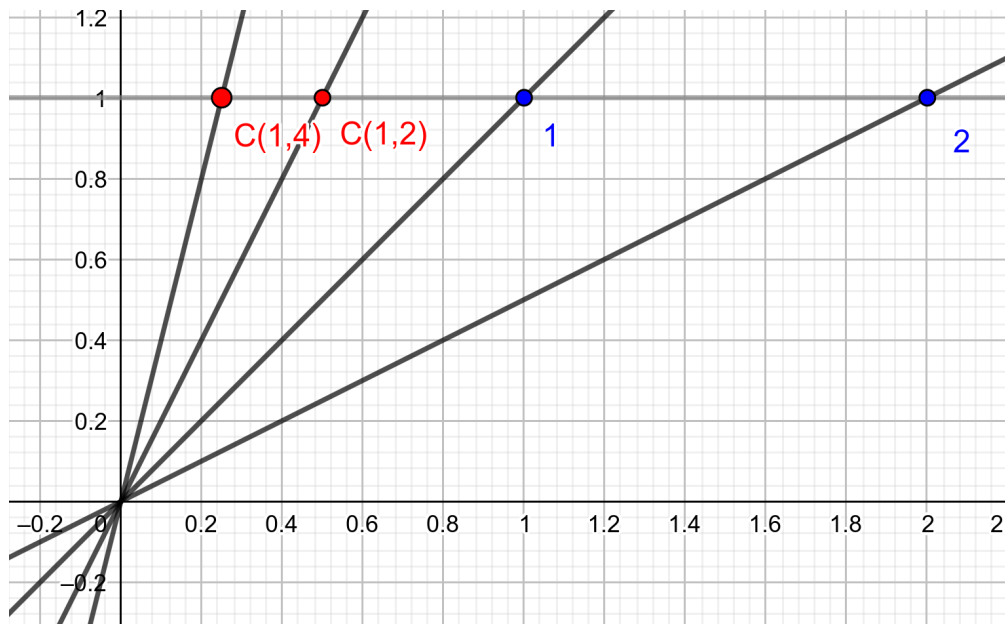


Abbildung 3.1: Der neue Zahlenstrahl ist die Gerade $f(x) = 1$. Die Schnittpunkte mit den anderen Geraden stellen die rationalen Zahlen dar.

In Proposition 2.2.5 haben wir gezeigt, dass jede Äquivalenzklasse in den ganzen Zahlen eine eindeutige „herausragende“ Darstellung besitzt. Das ist auch in den rationalen Zahlen der Fall.

Proposition 3.1.5. Sei $\frac{a}{b} \in \mathbb{Q}$ mit $a \neq 0$. Dann gibt es ein eindeutiges $(m, n) \in \mathbb{Z} \times (\mathbb{N} \setminus \{0\})$, sodass $\text{ggT}(m, n) = 1$ und $\frac{a}{b} = \frac{m}{n}$ ist.

Beweis. Sei $k = \text{ggT}(a, b)$. Da $k|a$ und $k|b$, gibt es ein $m \in \mathbb{Z}$ und ein $n \in \mathbb{N}$, sodass $a = k \cdot m$ und $b = k \cdot n$ ist. Es ist dann $\text{ggT}(m, n) = 1$, denn hätten m und n einen größeren gemeinsamen Teiler c , so wäre $m = c \cdot d_1$ und $n = c \cdot d_2$. Setzen wir das in die ersten Gleichungen für a und b ein, so erhalten wir $a = k \cdot c \cdot d_1$ und $b = k \cdot c \cdot d_2$. Es wäre also $k \cdot c$ ein gemeinsamer Teiler von a und b mit $k < k \cdot c$ – ein Widerspruch zur Maximalität von k . Somit muss $\text{ggT}(m, n) = 1$ sein.

Aus Proposition 3.1.4 folgt nun, dass $\frac{a}{b} = \frac{k \cdot m}{k \cdot n} = \frac{m}{n}$ ist.

Zur Eindeutigkeit: Angenommen, es gäbe $(\hat{m}, \hat{n}) \in \mathbb{Z} \times (\mathbb{N} \setminus \{0\})$ mit $\frac{a}{b} = \frac{m}{n} = \frac{\hat{m}}{\hat{n}}$ und $\text{ggT}(\hat{m}, \hat{n}) = 1$. Nach Definition 3.1.1 wäre damit $m \cdot \hat{n} = \hat{m} \cdot n$, es ist also m ein Teiler

von $n \cdot \hat{m}$. Da $\text{ggT}(m, n) = 1$ ist, muss nach Lemma 2.5.7 $m|\hat{m}$ gelten. Analog lässt sich zeigen, dass $\hat{m}|m$, $n|\hat{n}$ und $\hat{n}|n$ gelten. Daraus folgt mit Proposition 2.5.2, dass $n = \hat{n}$ und $|m| = |\hat{m}|$ ist. Hätten m und $\hat{m}$ verschiedene Vorzeichen, so wäre $m \cdot \hat{n} < n \cdot \hat{m}$ oder $m \cdot \hat{n} > n \cdot \hat{m}$ – ein Widerspruch zur Gleichheit. Somit ist auch $m = \hat{m}$. $\square$

3.2 Rechenoperationen auf $\mathbb{Q}$

Wir wollen in diesem Abschnitt die bekannten Rechenoperationen auf den rationalen Zahlen sauber definieren. Ziel wird sein, dass wir insgesamt eine Körperstruktur erhalten.

Für die Addition verwenden wir die aus der Schule bekannte Rechenregel. Wir zeigen, dass diese wohldefiniert ist und gemeinsam mit der noch zu definierenden Multiplikation einen Körper bildet.

Definition 3.2.1. Die Addition $+$: $\mathbb{Q} \times \mathbb{Q} \rightarrow \mathbb{Q}$ ist definiert durch

$$\frac{a}{b} + \frac{c}{d} := \frac{a \cdot d + c \cdot b}{b \cdot d}.$$

Proposition 3.2.2. Die Addition auf den rationalen Zahlen ist wohldefiniert. Ist also $(a', b') \in \frac{a}{b}$ und $(c', d') \in \frac{c}{d}$, so ist $\frac{a}{b} + \frac{c}{d} = \frac{a'}{b'} + \frac{c'}{d'}$.

Beweis. Es sind $(a, b) \sim (a', b')$ und $(c, d) \sim (c', d')$. Nach Definition 3.1.1 gilt daher $a \cdot b' = a' \cdot b$ und $c \cdot d' = c' \cdot d$. Multiplizieren wir die erste Gleichung mit $d \cdot d'$, die zweite Gleichung mit $b \cdot b'$, und addieren diese anschließend, so erhalten wir

$$a \cdot b' \cdot d \cdot d' + c \cdot d' \cdot b \cdot b' = a' \cdot b \cdot d \cdot d' + c' \cdot d \cdot b \cdot b'.$$

Durch Herausheben von Faktoren lässt sich diese Gleichung vereinfachen zu

$$(a \cdot d + c \cdot b) \cdot b' \cdot d' = (a' \cdot d' + c' \cdot b') \cdot b \cdot d.$$

Das bedeutet allerdings, dass $\frac{a \cdot d + c \cdot b}{b \cdot d} = \frac{a' \cdot d' + c' \cdot b'}{b' \cdot d'}$ ist, also $\frac{a}{b} + \frac{c}{d} = \frac{a'}{b'} + \frac{c'}{d'}$. $\square$

Sind $\frac{a}{b}, \frac{c}{b} \in \mathbb{Q}$ mit gleichem Nenner b , dann ist

$$\frac{a}{b} + \frac{c}{b} = \frac{a \cdot b + c \cdot b}{b \cdot b} = \frac{(a + c) \cdot b}{b \cdot b} = \frac{a + c}{b},$$

wobei für den letzten Schritt Proposition 3.1.4 verwendet wurde. Für das Addieren von gleichnamigen Brüchen müssen daher nur die Zähler addiert werden.

Definition 3.2.3. Die Multiplikation $\cdot : \mathbb{Q} \times \mathbb{Q} \rightarrow \mathbb{Q}$ ist definiert durch

$$\frac{a}{b} \cdot \frac{c}{d} := \frac{a \cdot c}{b \cdot d}.$$

Proposition 3.2.4. Die Multiplikation auf den rationalen Zahlen ist wohldefiniert. Ist also $(a', b') \in \frac{a}{b}$ und $(c', d') \in \frac{c}{d}$, so ist $\frac{a}{b} \cdot \frac{c}{d} = \frac{a'}{b'} \cdot \frac{c'}{d'}$.

Beweis. Es ist $(a, b) \sim (a', b')$ und $(c, d) \sim (c', d')$. Nach Definition 3.1.1 gilt daher $a \cdot b' = a' \cdot b$ und $c \cdot d' = c' \cdot d$. Multiplizieren wir diese Gleichungen miteinander und ordnen die Faktoren um, so erhalten wir

$$a \cdot c \cdot b' \cdot d' = a' \cdot c' \cdot b \cdot d.$$

Es ist somit $\frac{a \cdot c}{b \cdot d} = \frac{a' \cdot c'}{b' \cdot d'}$ nach Definition 3.1.1. □

Proposition 3.2.5. $(\mathbb{Q}, +, \cdot, \frac{0}{1}, \frac{1}{1})$ ist ein Körper. Es sind daher die Körperaxiome erfüllt (vgl. [50, S. 264]):

- (1) $\forall a, b, c \in \mathbb{Q} : (a + b) + c = a + (b + c)$ (Assoziativität von $+$).
- (2) $\forall a, b \in \mathbb{Q} : a + b = b + a$ (Kommutativität von $+$).
- (3) $\forall a \in \mathbb{Q} : a + \frac{0}{1} = a$ (Nullelement).
- (4) $\forall a \in \mathbb{Q} : \exists -a \in \mathbb{Q} : a + (-a) = \frac{0}{1}$ (Invertierbarkeit von $+$).
- (5) $\forall a, b, c \in \mathbb{Q} : (a \cdot b) \cdot c = a \cdot (b \cdot c)$ (Assoziativität von $\cdot$).
- (6) $\forall a, b \in \mathbb{Q} : a \cdot b = b \cdot a$ (Kommutativität von $\cdot$).
- (7) $\forall a \in \mathbb{Q} : a \cdot \frac{1}{1} = a$ (Einselement).

$$(8) \quad \forall a \in \mathbb{Q} \setminus \left\{ \frac{0}{1} \right\} : \exists a^{-1} \in \mathbb{Q} : a \cdot a^{-1} = \frac{1}{1} \text{ (Invertierbarkeit von } \cdot \text{)}.$$

$$(9) \quad \forall a, b, c \in \mathbb{Q} : a \cdot (b + c) = a \cdot b + a \cdot c \text{ (Distributivität)}.$$

Beweis. Seien $a = \frac{a_1}{a_2}$, $b = \frac{b_1}{b_2}$ und $c = \frac{c_1}{c_2}$ rationale Zahlen. Wir beweisen zunächst, dass das Assoziativgesetz der Addition erfüllt ist:

$$\begin{aligned} (a + b) + c &= \left(\frac{a_1}{a_2} + \frac{b_1}{b_2} \right) + \frac{c_1}{c_2} = \\ &= \frac{a_1 \cdot b_2 + b_1 \cdot a_2}{a_2 \cdot b_2} + \frac{c_1}{c_2} = \\ &= \frac{a_1 \cdot b_2 \cdot c_2 + b_1 \cdot a_2 \cdot c_2 + c_1 \cdot a_2 \cdot b_2}{a_2 \cdot b_2 \cdot c_2} = \\ &= \frac{a_1 \cdot b_2 \cdot c_2 + (b_1 \cdot c_2 + c_1 \cdot b_2) \cdot a_2}{a_2 \cdot b_2 \cdot c_2} = \\ &= \frac{a_1}{a_2} + \frac{b_1 \cdot c_2 + c_1 \cdot b_2}{b_2 \cdot c_2} = \\ &= \frac{a_1}{a_2} + \left(\frac{b_1}{b_2} + \frac{c_1}{c_2} \right) = \\ &= a + (b + c). \end{aligned}$$

Die Kommutativität der Addition folgt direkt aus der Kommutativität der Addition und Multiplikation in den ganzen Zahlen:

$$a + b = \frac{a_1}{a_2} + \frac{b_1}{b_2} = \frac{a_1 \cdot b_2 + b_1 \cdot a_2}{a_2 \cdot b_2} = \frac{b_1 \cdot a_2 + a_1 \cdot b_2}{b_2 \cdot a_2} = \frac{b_1}{b_2} + \frac{a_1}{a_2} = b + a.$$

Durch Einsetzen in die Definition folgt, dass es sich bei $\frac{0}{1}$ um das Nullelement der Addition handelt:

$$a + \frac{0}{1} = \frac{a_1}{a_2} + \frac{0}{1} = \frac{a_1 \cdot 1 + 0 \cdot a_2}{a_2 \cdot 1} = \frac{a_1}{a_2} = a.$$

Um die Invertierbarkeit der Addition zu beweisen, definieren wir zunächst

$$-a := \frac{-a_1}{a_2}.$$

Es ist dann

$$a + (-a) = \frac{a_1}{a_2} + \frac{-a_1}{a_2} = \frac{a_1 \cdot a_2 + (-a_1) \cdot a_2}{a_2 \cdot a_2} = \frac{0}{a_2^2} = \frac{0}{1},$$

wobei wir für die letzte Gleichheit Proposition 3.1.4 verwendet haben.

Die Assoziativität der Multiplikation lässt sich beweisen, indem wir wieder die entsprechende Eigenschaft der Multiplikation in den ganzen Zahlen verwenden:

$$\begin{aligned} (a \cdot b) \cdot c &= \left(\frac{a_1}{a_2} \cdot \frac{b_1}{b_2} \right) \cdot \frac{c_1}{c_2} = \\ &= \frac{a_1 \cdot b_1}{a_2 \cdot b_2} \cdot \frac{c_1}{c_2} = \\ &= \frac{(a_1 \cdot b_1) \cdot c_1}{(a_2 \cdot b_2) \cdot c_2} = \\ &= \frac{a_1 \cdot (b_1 \cdot c_1)}{a_2 \cdot (b_2 \cdot c_2)} = \\ &= \frac{a_1}{a_2} \cdot \frac{b_1 \cdot c_1}{b_2 \cdot c_2} = \\ &= \frac{a_1}{a_2} \cdot \left(\frac{b_1}{b_2} \cdot \frac{c_1}{c_2} \right) = \\ &= a \cdot (b \cdot c). \end{aligned}$$

Die Kommutativität der Multiplikation folgt analog:

$$a \cdot b = \frac{a_1}{a_2} \cdot \frac{b_1}{b_2} = \frac{a_1 \cdot b_1}{a_2 \cdot b_2} = \frac{b_1 \cdot a_1}{b_2 \cdot a_2} = \frac{b_1}{b_2} \cdot \frac{a_1}{a_2} = b \cdot a.$$

Einselement der Multiplikation:

$$a \cdot \frac{1}{1} = \frac{a_1}{a_2} \cdot \frac{1}{1} = \frac{a_1 \cdot 1}{a_2 \cdot 1} = \frac{a_1}{a_2} = a.$$

Wir wollen nun zeigen, dass jede rationale Zahl ungleich dem Nullelement ein multiplikatives Inverses besitzt. Sei dazu $a \neq \frac{0}{1}$. Wir setzen

$$a^{-1} := \frac{a_2}{a_1}.$$

Dann ist

$$a \cdot a^{-1} = \frac{a_1}{a_2} \cdot \frac{a_2}{a_1} = \frac{a_1 \cdot a_2}{a_2 \cdot a_1} = \frac{1}{1},$$

wobei wir bei der letzten Umformung mit Hilfe von Proposition 3.1.4 gekürzt haben.

Zuletzt beweisen wir die Distributivität:

$$\begin{aligned} a \cdot (b + c) &= \frac{a_1}{a_2} \cdot \left(\frac{b_1}{b_2} + \frac{c_1}{c_2} \right) = \\ &= \frac{a_1}{a_2} \cdot \frac{b_1 \cdot c_2 + c_1 \cdot b_2}{b_2 \cdot c_2} = \\ &= \frac{a_1 \cdot (b_1 \cdot c_2 + c_1 \cdot b_2)}{a_2 \cdot b_2 \cdot c_2} = \\ &= \frac{a_1 \cdot b_1 \cdot c_2 + a_1 \cdot c_1 \cdot b_2}{a_2 \cdot b_2 \cdot c_2} = \\ &= \frac{a_1 \cdot b_1 \cdot c_2}{a_2 \cdot b_2 \cdot c_2} + \frac{a_1 \cdot c_1 \cdot b_2}{a_2 \cdot b_2 \cdot c_2} = \\ &= \frac{a_1 \cdot b_1}{a_2 \cdot b_2} + \frac{a_1 \cdot c_1}{a_2 \cdot c_2} = \\ &= \frac{a_1}{a_2} \cdot \frac{b_1}{b_2} + \frac{a_1}{a_2} \cdot \frac{c_1}{c_2} = \\ &= a \cdot b + a \cdot c. \end{aligned}$$

□

Für das additive Inverse zu $a \in \mathbb{Q}$ vereinbaren wir wieder die Schreibweise $-a$, während wir das multiplikative Inverse von a als a^{-1} bezeichnen. Zudem schreiben wir anstelle von $a + (-b)$ direkt $a - b$.

Wir zeigen nun einige Resultate, die allgemein für Körper, und daher speziell für die rationalen Zahlen, gelten.

Proposition 3.2.6. *Sei $(K, +, \cdot, 0, 1)$ ein Körper und $a, b, c \in K$. Dann gelten folgende Eigenschaften:*

- (1) *K ist nullteilerfrei.*
- (2) *Es gelten die Kürzungsregeln: Aus $a + c = b + c$ folgt $a = b$, und ist $c \neq 0$, so folgt aus $a \cdot c = b \cdot c$ ebenfalls, dass $a = b$ ist.*
- (3) *Die Gleichung $a \cdot x + b = c$ mit $a \neq 0$ besitzt stets eine eindeutige Lösung $x \in K$.*

Beweis. Zur Nullteilerfreiheit: Angenommen, es wäre $a \cdot b = 0$. Ist $a = 0$, so sind wir fertig. Andernfalls gibt es zu a ein multiplikatives Inverses a^{-1} . Multiplizieren der Gleichung mit a^{-1} liefert

$$0 = a^{-1} \cdot a \cdot b = 1 \cdot b = b.$$

Die Kürzungsregeln folgen direkt durch Addieren von $-c$, beziehungsweise durch Multiplizieren mit c^{-1} auf beiden Seiten der Gleichung.

Für den dritten Punkt betrachten wir $x := a^{-1} \cdot (c - b)$. Durch Einsetzen wird sofort klar, dass es sich bei x um eine Lösung der Gleichung handelt. Gäbe es noch eine zweite Lösung x' , so wäre $a \cdot x + b = c$ und $a \cdot x' + b = c$. Subtrahieren der zweiten Gleichung von der ersten liefert $a \cdot (x - x') = 0$. Da $a \neq 0$ ist, und ein Körper wegen des ersten Punkts nullteilerfrei ist, muss $x = x'$ sein. $\square$

Wir wollen nun eine Ordnung auf den rationalen Zahlen einführen. In den ganzen Zahlen haben wir definiert, dass für $a, b \in \mathbb{Z}$ die Ungleichung $a \leq b$ gilt, falls es eine nichtnegative ganze Zahl k gibt, sodass $b = a + k$ ist. Diese Definition wollen wir auf die rationalen Zahlen übertragen, wofür wir zunächst die Eigenschaften „positiv“ und „negativ“ auf den rationalen Zahlen definieren müssen.

Definition 3.2.7. Sei $\frac{a}{b} \in \mathbb{Q}$ mit $\frac{a}{b} \neq \frac{0}{1}$, und $\frac{m}{n} = \frac{a}{b}$ die nach Proposition 3.1.5 existierende eindeutige Darstellung von $\frac{a}{b}$ mit $\text{ggT}(m, n) = 1$. Wir nennen $\frac{a}{b}$ positiv, falls m positiv ist, und nennen den Bruch negativ, falls m negativ ist. Das Nullelement $\frac{0}{1}$ wird als weder positiv noch negativ definiert.

Zudem definieren wir

$$\begin{aligned} \mathbb{Q}^+ &:= \left\{ \frac{m}{n} \in \mathbb{Q} \mid \text{ggT}(m, n) = 1 \wedge m \in \mathbb{Z}^+ \right\}, \\ \mathbb{Q}^- &:= \left\{ \frac{m}{n} \in \mathbb{Q} \mid \text{ggT}(m, n) = 1 \wedge m \in \mathbb{Z}^- \right\} \end{aligned}$$

als die Menge der positiven, beziehungsweise der negativen rationalen Zahlen.

Ist eine rationale Zahl positiv, so ist ihr Zähler ebenfalls stets positiv – unabhängig von der Darstellung der rationalen Zahl. Analoges gilt für negative rationale Zahlen, und wird durch folgende Proposition formalisiert.

Proposition 3.2.8. *Sei $\frac{a}{b} \in \mathbb{Q}$. Dann ist $\frac{a}{b} \in \mathbb{Q}^+$ genau dann, wenn $a \in \mathbb{Z}^+$. Weiters ist $\frac{a}{b} \in \mathbb{Q}^-$ genau dann, wenn $a \in \mathbb{Z}^-$, und $\frac{a}{b} = \frac{0}{1}$ äquivalent dazu, dass $a = 0$ ist.*

Beweis. Wir beginnen mit dem Fall $\frac{a}{b} = \frac{0}{1}$. Nach Definition 3.1.1 ist dann $a \cdot 1 = 0 \cdot b = 0$, also $a = 0$. Umgekehrt folgt aus $a = 0$ mit Proposition 3.1.4, dass $\frac{0 \cdot b}{1 \cdot b} = \frac{0}{1}$ ist.

Angenommen, $\frac{a}{b}$ ist positiv und $\frac{m}{n} = \frac{a}{b}$ mit $\text{ggT}(m, n) = 1$. Nach Definition 3.1.1 ist das äquivalent zu $m \cdot b = a \cdot n$. Da beide Seiten der Gleichung dasselbe Vorzeichen haben müssen, ist daher m genau dann positiv, wenn es a ist.

Falls $\frac{a}{b}$ negativ ist, so muss nach Definition $m \in \mathbb{Z}^-$ sein. Wieder müssen in der Gleichung $m \cdot b = a \cdot n$ beide Seiten dasselbe Vorzeichen besitzen, womit m genau dann negativ ist, wenn es a ist. $\square$

Proposition 3.2.9. *Es ist $\mathbb{Q} = \mathbb{Q}^- \cup \left\{\frac{0}{1}\right\} \cup \mathbb{Q}^+$ und $\mathbb{Q}^- \cap \mathbb{Q}^+ = \emptyset$. Anders ausgedrückt: Jede rationale Zahl ist entweder positiv oder negativ, oder gleich dem Nullelement.*

Beweis. Es ist offensichtlich $\frac{0}{1}$ Element dieser Vereinigung. Für ein $\frac{a}{b}$ in $\mathbb{Q} \setminus \left\{\frac{0}{1}\right\}$ gibt es nach Proposition 3.1.5 eine Darstellung $\frac{m}{n} = \frac{a}{b}$ mit $\text{ggT}(m, n) = 1$. Da $m \in \mathbb{Z}$ ist, und jede ganze Zahl ungleich Null nach Proposition 2.3.12 entweder positiv oder negativ ist, muss $\frac{a}{b}$ entweder in $\mathbb{Q}^-$ oder in $\mathbb{Q}^+$ liegen. Da es keine ganze Zahl gibt, die positiv und negativ ist, folgt daraus insbesondere, dass $\mathbb{Q}^- \cap \mathbb{Q}^+ = \emptyset$ ist. $\square$

Die aus den ganzen Zahlen bekannten Rechenregeln (siehe Abschnitt 2.3) gelten analog in den rationalen Zahlen. Denn ist zum Beispiel $\frac{a}{b} \in \mathbb{Q}^-$ und $\frac{c}{d} \in \mathbb{Q}^+$, so ist $\frac{a}{b} \cdot \frac{c}{d} = \frac{a \cdot c}{b \cdot d}$, wobei $a \cdot c$ negativ ist. Das Ergebnis ist daher eine negative rationale Zahl.

Lemma 3.2.10. *Sei $\frac{a}{b} \in \mathbb{Q}$. Falls $\frac{a}{b}$ negativ ist, so ist $-\frac{a}{b}$ positiv. Ist hingegen $\frac{a}{b}$ positiv, so ist $-\frac{a}{b}$ negativ.*

Beweis. Ist $\frac{a}{b}$ negativ, so ist auch a negativ, was nach Lemma 2.3.13 bedeutet, dass $-a$ positiv ist. Mit Proposition 3.2.8 folgt direkt, dass $\frac{-a}{b} = -\frac{a}{b}$ positiv ist.

Die zweite Behauptung lässt sich analog beweisen: Ist $\frac{a}{b}$ positiv, so ist a positiv, und damit $-a$ negativ. Insgesamt ist daher $\frac{-a}{b} = -\frac{a}{b}$ negativ. $\square$

Definition 3.2.11. Seien $a, b \in \mathbb{Q}$. Dann definieren wir

$$a \leq b : \Leftrightarrow \exists c \in \mathbb{Q} \setminus \mathbb{Q}^- : a + c = b.$$

Wir schreiben zudem $b \geq a$ für $a \leq b$. Ist das obige c positiv (also ungleich $\frac{0}{1}$), so verwenden wir $<$ und $>$ anstelle von $\leq$ und $\geq$.

Wenig überraschend ist, dass die negativen rationalen Zahlen kleiner als $\frac{0}{1}$ sind, während die positiven rationalen Zahlen größer als $\frac{0}{1}$ sind. Denn für eine negative rationale Zahl $\frac{a}{b}$ ist $-a$ nach Lemma 2.3.13 positiv und $\frac{a}{b} + \frac{-a}{b} = \frac{0}{b} = \frac{0}{1}$. Ist $\frac{a}{b}$ hingegen positiv, so ist offensichtlich $\frac{0}{1} + \frac{a}{b} = \frac{a}{b}$, was $\frac{0}{1} < \frac{a}{b}$ bedeutet.

Proposition 3.2.12. Die Relation $\leq$ auf $\mathbb{Q}$ ist eine Totalordnung. Es gilt daher $a \leq b$ oder $b \leq a$ für alle $a, b \in \mathbb{Q}$. Zudem werden die Eigenschaften einer Ordnungsrelation erfüllt:

- (1) $\forall a \in \mathbb{Q} : a \leq a$ (Reflexivität).
- (2) $\forall a, b \in \mathbb{Q} : (a \leq b \wedge b \leq a) \Rightarrow a = b$ (Antisymmetrie).
- (3) $\forall a, b, c \in \mathbb{Q} : (a \leq b \wedge b \leq c) \Rightarrow a \leq c$ (Transitivität).

Beweis. Zur Reflexivität: Es ist offensichtlich $a + \frac{0}{1} = a$, und damit $a \leq a$.

Zur Antisymmetrie: Angenommen, für $a, b \in \mathbb{Q}$ ist $a \leq b$ und $b \leq a$. Es gibt daher $c, c' \in \mathbb{Q} \setminus \mathbb{Q}^-$, sodass $b = a + c$ und $a = b + c'$ ist. Einsetzen der ersten Gleichung in die zweite liefert $a = a + (c + c')$, also $c + c' = \frac{0}{1}$. Wäre c oder c' positiv, so wäre auch die Summe positiv – im Widerspruch dazu, dass $\frac{0}{1}$ weder positiv noch negativ ist. Es kann somit nur $c = c' = \frac{0}{1}$ sein, woraus $a = b$ folgt.

Zur Transitivität: Falls $a \leq b$ und $b \leq c$ ist, so gibt es $d, d' \in \mathbb{Q} \setminus \mathbb{Q}^-$, sodass $b = a + d$ und $c = b + d'$ ist. Einsetzen der ersten Gleichung in die zweite liefert $c = a + (d + d')$,

wobei $d + d'$ als Summe zweier nichtnegativer rationaler Zahlen wieder nichtnegativ ist. Daher ist auch $a \leq c$.

Um zu zeigen, dass stets zwei rationale Zahlen vergleichbar sind, betrachten wir $a, b \in \mathbb{Q}$. Nach Proposition 3.2.6 besitzt die Gleichung $a + x = b$ stets eine eindeutige Lösung $x \in \mathbb{Q}$. Ist diese nichtnegativ, so gilt $a \leq b$. Falls x hingegen negativ ist, so ist $-x$ positiv, und $a = b + (-x)$, was $b \leq a$ bedeutet. $\square$

Proposition 3.2.13. *Die Relation $\leq$ auf den rationalen Zahlen erfüllt die Ordnungsaxiome:*

(1) Für $a, b, c \in \mathbb{Q}$ ist $a \leq b \Leftrightarrow a + c \leq b + c$.

(2) Für $a, b \in \mathbb{Q}$ mit $a > \frac{0}{1}$ und $b > \frac{0}{1}$ ist $a \cdot b > \frac{0}{1}$.

Beweis. Zu (1): Es ist $a \leq b$ äquivalent dazu, dass es ein $d \in \mathbb{Q} \setminus \mathbb{Q}^-$ gibt, sodass $b = a + d$ ist. Diese Gleichung ist wiederum äquivalent zu $b + c = a + c + d$, was $a + c \leq b + c$ bedeutet.

Zu (2): Nach den vorigen Beobachtungen ist das Produkt zweier positiver rationaler Zahlen wieder positiv. Da die positiven rationalen Zahlen größer als $\frac{0}{1}$ sind, folgt damit $a \cdot b > \frac{0}{1}$. $\square$

Mit Ungleichungen in den rationalen Zahlen lässt sich weiterhin wie gewohnt rechnen. Die aus den ganzen Zahlen hinzugekommene Eigenschaft, dass sich beim Multiplizieren einer Ungleichung mit einer negativen Zahl das Ungleichheitszeichen umdreht, bleibt erhalten.

Proposition 3.2.14. *Seien $a, b, c, d \in \mathbb{Q}$. Dann gilt:*

(1) Ist $a \leq b$ und $c \leq d$, so ist $a + c \leq b + d$.

(2) Ist c positiv, so ist $a \leq b$ äquivalent zu $a \cdot c \leq b \cdot c$.

(3) Ist c negativ, so ist $a \leq b$ äquivalent zu $b \cdot c \leq a \cdot c$.

Beweis. Zu (1): Da $a \leq b$ und $c \leq d$ ist, gibt es $q, q' \in \mathbb{Q} \setminus \mathbb{Q}^-$, sodass $b = a + q$ und $d = c + q'$ ist. Addieren dieser Gleichungen liefert $b + d = a + c + (q + q')$, was nach Definition 3.2.11 bedeutet, dass $a + c \leq b + c$ ist.

Zu (2): Ist $a \leq b$, so gibt es wieder ein $q \in \mathbb{Q} \setminus \mathbb{Q}^-$ mit $b = a + q$. Multiplizieren dieser Gleichung mit c und Anwenden des Distributivgesetzes liefert $b \cdot c = a \cdot c + q \cdot c$. Da sowohl q als auch c nichtnegativ waren, ist auch das Produkt $q \cdot c$ nichtnegativ. Damit ist $a \cdot c \leq b \cdot c$.

Ist umgekehrt $a \cdot c \leq b \cdot c$, so gibt es ein $q \in \mathbb{Q} \setminus \mathbb{Q}^-$, sodass $b \cdot c = a \cdot c + q$ ist. Multiplizieren dieser Gleichung mit c^{-1} liefert $b = a + q \cdot c^{-1}$. Da c^{-1} nur positiv sein kann (andernfalls wäre $\frac{1}{c} = c \cdot c^{-1}$ negativ), ist auch $q \cdot c^{-1}$ nichtnegativ, was $a \leq b$ bedeutet.

Zu (3): Ist c negativ, so ist nach Lemma 3.2.10 $-c$ positiv. Nach (2) ist damit $a \leq b$ äquivalent zu $a \cdot (-c) \leq b \cdot (-c)$. Es gibt daher ein $q \in \mathbb{Q} \setminus \mathbb{Q}^-$, sodass $b \cdot (-c) = a \cdot (-c) + q$ ist. Multiplizieren dieser Gleichung mit $-\frac{1}{c}$ liefert $b \cdot c = a \cdot c - q$. Addieren wir nun auf beiden Seiten der Gleichung q , so erhalten wir $a \cdot c = b \cdot c + q$, was nach Definition 3.2.11 $b \cdot c \leq a \cdot c$ bedeutet. $\square$

Die Betragsfunktion der ganzen Zahlen lässt sich auf die rationalen Zahlen erweitern. Wieder sind alle bekannten Eigenschaften erfüllt.

Definition 3.2.15. Sei $a \in \mathbb{Q}$. Dann definieren wir

$$|a| := \begin{cases} a & \text{wenn } a \geq \frac{0}{1}, \\ -a & \text{sonst.} \end{cases}$$

Die Ordnung aus Definition 3.2.11 stellt, wie schon in den ganzen Zahlen, keine Wohlordnung dar. Denn wäre $\leq$ eine Wohlordnung auf den rationalen Zahlen, so müsste auch die Menge der negativen rationalen Zahlen $\mathbb{Q}^-$ ein kleinstes Element besitzen – nennen wir dieses $\frac{a}{b}$. Es wäre dann $\frac{a}{b} + \frac{-b}{b} = \frac{a+(-b)}{b}$ eine negative ganze Zahl, denn die Summe $a + (-b)$ zweier negativer ganzer Zahlen ist wieder negativ. Damit ist $\frac{a}{b} = \frac{a+(-b)}{b} + \frac{b}{b}$, also $\frac{a+(-b)}{b} < \frac{a}{b}$, ein Widerspruch zur Minimalität von $\frac{a}{b}$.

3.3 Einbettung von $\mathbb{Z}$ in $\mathbb{Q}$

Wir wollen nun, wie im vorigen Kapitel, unsere bisherige Notation der ganzen Zahlen auf die rationalen Zahlen übertragen. Gesucht ist daher eine Einbettung von $\mathbb{Z}$ in $\mathbb{Q}$.

Intuitiv fällt uns die Wahl nicht sonderlich schwer: Wir „wissen“ aus der Schule, dass wir jede ganze Zahl m als $\frac{m}{1}$ anschreiben können. Folgende Proposition rechtfertigt das.

Proposition 3.3.1. *Sei $\varphi : \mathbb{Z} \rightarrow \mathbb{Q}$ mit $\varphi(m) := \frac{m}{1}$. Dann handelt es sich bei φ um eine Einbettung.*

Beweis. Wir müssen die geforderten drei Eigenschaften für eine Einbettung nachweisen. Seien dazu $n, m \in \mathbb{Z}$. Dann ist $\varphi(n + m) = \frac{n+m}{1} = \frac{n}{1} + \frac{m}{1} = \varphi(m) + \varphi(n)$. Es ist zudem $\varphi(n \cdot m) = \frac{n \cdot m}{1} = \frac{n}{1} \cdot \frac{m}{1} = \varphi(n) \cdot \varphi(m)$, und somit auch die zweite Eigenschaft erfüllt.

Für den Beweis der Injektivität von φ nehmen wir an, es wäre $\varphi(m) = \frac{m}{1} = \frac{n}{1} = \varphi(n)$. Nach Definition 3.1.1 ist damit $m \cdot 1 = 1 \cdot m$, also $m = n$. Die Funktion φ ist daher injektiv, und somit eine Einbettung. $\square$

Die Ordnung aus den ganzen Zahlen überträgt sich auf die rationalen Zahlen, wie folgende Proposition zeigt.

Proposition 3.3.2. *Sei φ die Einbettung aus Proposition 3.3.1. Dann ist für $n, m \in \mathbb{Z}$ $n \leq m \Leftrightarrow \varphi(n) \leq \varphi(m)$.*

Beweis. Falls $n \leq m$, so gibt es ein $k \in \mathbb{Z} \setminus \mathbb{Z}^-$, sodass $m = n + k$ ist. Es ist dann $\varphi(m) = \varphi(n + k) = \varphi(n) + \varphi(k)$, wobei $\varphi(k) \in \mathbb{Q} \setminus \mathbb{Q}^-$. Damit ist $\varphi(n) \leq \varphi(m)$.

Gilt umgekehrt $\varphi(n) \leq \varphi(m)$, so gibt es ein $\frac{a}{b} \in \mathbb{Q} \setminus \mathbb{Q}^-$ mit $\text{ggT}(a, b) = 1$, sodass $\varphi(m) = \varphi(n) + \frac{a}{b}$ ist. Es ist daher

$$\frac{m}{1} = \frac{n}{1} + \frac{a}{b} = \frac{n \cdot b + a}{b},$$

und damit $m \cdot b = n \cdot b + a$ wegen Definition 3.1.1, was sich zu $(m - n) \cdot b = a$ umformen lässt. Daher ist $b|a$ und zugleich $\text{ggT}(a, b) = 1$, also $b = 1$. Damit ist $\frac{m}{1} = \frac{n}{1} + \frac{a}{1}$, also $\varphi(m) = \varphi(n) + \varphi(a)$, wobei a nach Proposition 3.2.8 nichtnegativ ist. Somit ist $n \leq m$. $\square$

Die ganzen Zahlen sind daher mit den rationalen Zahlen verträglich. Für die übliche Notation müssen wir nur wenige Vereinbarungen treffen: Lässt sich ein Bruch auf die Gestalt $\frac{m}{1}$ bringen, so schreiben wir diesen künftig auch als m an; umgekehrt können wir jede ganze Zahl m als $\frac{m}{1} = \frac{2 \cdot m}{2} = \dots$ anschreiben [34, S. 19].

Wie schon zuvor schreiben wir $\mathbb{Z} \subseteq \mathbb{Q}$ um auszudrücken, dass die ganzen Zahlen in den rationalen Zahlen eingebettet sind.

Gegenüber den ganzen Zahlen haben wir in den rationalen Zahlen die Eigenschaft gewonnen, dass für jede Zahl ungleich Null ein multiplikatives Inverses existiert. Obwohl uns $\mathbb{Q}$ als Körper viele schöne Eigenschaften garantiert, werden wir im nächsten Kapitel feststellen, dass der Zahlenstrahl in Abbildung 3.1 nur eine durchgehende Linie vortäuscht – tatsächlich sind darin unzählig viele Löcher enthalten.

Eine weitere Möglichkeit, rationale Zahlen darzustellen, ist durch Dezimalzahlen gegeben (etwa $\frac{1}{8} = 0,125$). Während wir bei den Zahlen vor dem Komma im Stellenwertsystem nach Einern, Zehnern, Hundertern, sprich nach Potenzen von 10 zusammenfassen, bündeln wir bei den Nachkommastellen von links ausgehend nach Zehnteln, Hundertsteln, Tausendsteln, also nach Potenzen von $\frac{1}{10}$.

Proposition 3.3.3. Sei $\frac{a}{b} \in \mathbb{Q}$. Dann gibt es eine Folge ganzer Zahlen $(c_n)_{n \in \mathbb{N}}$ mit $c_0 \in \mathbb{Z}$ und $0 \leq c_n \leq 9$ für alle $n \geq 1$, sodass

$$\frac{a}{b} = c_0 + \frac{c_1}{10} + \frac{c_2}{100} + \frac{c_3}{1000} + \dots = c_0 + \sum_{n=1}^{\infty} c_n \cdot \left(\frac{1}{10}\right)^n.$$

Wir schreiben dann $\frac{a}{b} = c_0, c_1 c_2 c_3 \dots$

Beweis. Nach Proposition 2.5.3 gibt es ein $c_0 \in \mathbb{Z}$ und ein $r_0 \in \mathbb{N}$ mit $0 \leq r_0 < b$, sodass $a = c_0 \cdot b + r_0$ ist. Für $n \in \mathbb{N}$ definieren wir c_{n+1} und r_{n+1} nun rekursiv:

$$10 \cdot r_n = c_{n+1} \cdot b + r_{n+1}, \text{ wobei } 0 \leq r_{n+1} < b.$$

Für $n \geq 1$ ist dann $c_n \in \mathbb{N}$, da sowohl r_n als auch b stets natürliche Zahlen sind. Wegen $r_n < b$ ist $0 \leq c_n \leq 9$.

Betrachten wir nun den folgenden Ausdruck:

$$c_0 \cdot b + \sum_{n=1}^{\infty} c_n \cdot b \cdot \left(\frac{1}{10}\right)^n.$$

Mithilfe von $10 \cdot r_{n-1} = c_n \cdot b + r_n$ lässt sich das umformen zu

$$c_0 \cdot b + \sum_{n=1}^{\infty} (10 \cdot r_{n-1} - r_n) \cdot \left(\frac{1}{10}\right)^n.$$

Bei dieser Reihe handelt es sich um eine Teleskopsumme, die gegen r_0 konvergiert. Es ist daher

$$a = c_0 \cdot b + r_0 = c_0 \cdot b + \sum_{n=1}^{\infty} c_n \cdot b \cdot \left(\frac{1}{10}\right)^n,$$

woraus wir durch Multiplikation mit $\frac{1}{b}$ direkt

$$\frac{a}{b} = c_0 + \sum_{n=1}^{\infty} c_n \cdot \left(\frac{1}{10}\right)^n$$

erhalten. □

Die Dezimaldarstellung ist nicht eindeutig, wie das bekannte Beispiel $1 = 0, \bar{9}$ zeigt:

$$0, \bar{9} = \sum_{n=1}^{\infty} 9 \cdot \left(\frac{1}{10}\right)^n = 9 \cdot \sum_{n=1}^{\infty} \left(\frac{1}{10}\right)^n = 9 \cdot \frac{1}{1-10} = \frac{9}{9} = 1.$$

Jede rationale Zahl besitzt eine endliche oder eine unendliche periodische Dezimaldarstellung.

Proposition 3.3.4. Sei $a = a_0, a_1 a_2 a_3 \dots$ eine Dezimaldarstellung einer rationalen Zahl a . Dann gilt einer der Fälle:

- (1) Es gibt ein $N \in \mathbb{N}$, sodass $a_n = 0$ für alle $n > N$ ist. Die Dezimaldarstellung bricht daher nach der N -ten Nachkommastelle ab.

- (2) Es gibt $N, M \in \mathbb{N} \setminus \{0\}$, sodass $a_n = a_{n-M}$ für alle $n > N$ ist, wobei ein $K > N$ mit $a_K \neq 0$ existiert. In dem Fall handelt es sich um eine periodische Dezimalzahl mit Periodenlänge M .

Umgekehrt handelt es sich bei jeder Dezimalzahl, die eine dieser Eigenschaften erfüllt, um eine rationale Zahl.

Beweis. Betrachten wir die Methode, mit der wir in Proposition 3.3.3 die Dezimaldarstellung einer rationalen Zahl erhalten haben. Falls wir durch Anwenden dieser Methode bei einer Division mit Rest jemals den Rest $r_N = 0$ erhalten, so ist $a_n = r_n = 0$ für alle $n > N$. Alle Nachkommastellen ab der $(N + 1)$ -ten Stelle sind daher gleich Null.

Erhalten wir hingegen bei keiner Division den Rest Null, so gilt jedenfalls, dass die Reste der Division stets die Ungleichung $0 < r_n < b$ für alle $n \in \mathbb{N}$ erfüllen. Da die Menge der möglichen Reste beschränkt ist, wir allerdings beliebig viele Reste haben, muss bei einem Index $N \in \mathbb{N}$ einer der Reste erstmals doppelt vorkommen. Es gibt daher ein $M \in \mathbb{N}$ mit $r_N = r_{N-M}$. Da $10 \cdot r_n = a_{n+1} \cdot b + r_{n+1}$ ist, muss wegen der Eindeutigkeit der Division $a_{N+1} = a_{N+1-M}$ sein, und damit $a_n = a_{n-M}$ für alle $n > N$.

Wäre $a_0, a_1 a_2 \dots a_N = a_0 + \sum_{n=1}^N a_n \cdot \left(\frac{1}{10}\right)^n$ eine abbrechende Dezimalzahl, so folgt direkt aus der Darstellung als endliche Summe, dass es sich um eine rationale Zahl handelt.

Handelt es sich hingegen um eine periodische Dezimalzahl

$$a = a_0, a_1 a_2 \dots \overline{a_N a_{N+1} \dots a_{N+(M-1)}}$$

mit Periodenlänge $M \geq 1$, so können wir uns zunächst überlegen, dass es ausreicht, die Rationalität von $0, \overline{a_N a_{N+1} \dots a_{N+(M-1)}}$ zu zeigen. Ist diese periodische Dezimalzahl rational, so ist sie es auch, wenn wir sie mit $\left(\frac{1}{10}\right)^{N-1}$ multiplizieren, wodurch wir den periodischen Dezimalteil von a erhalten. Addieren wir dazu $a_0, a_1 a_2 \dots a_{N-1}$ (was nach unserer vorigen Überlegung eine rationale Zahl ist), so erhalten wir die ursprüngliche Zahl a als eine Summe rationaler Zahlen. Damit ist a rational.

Wir wollen also zeigen, dass $b = 0, \overline{a_N a_{N+1} \dots a_{N+(M-1)}}$ rational ist. Dazu multiplizieren

wir diese Zahl mit 10^M und erhalten

$$10^M \cdot b = a_N a_{N+1} \dots a_{N+(M-1)}, \overline{a_N a_{N+1} \dots a_{N+(M-1)}}.$$

Subtrahieren wir b , bekommen wir

$$10^M \cdot b - b = b \cdot (10^M - 1) = a_N a_{N+1} \dots a_{N+(M-1)} \in \mathbb{Z}.$$

Daher ist

$$b = \frac{a_N a_{N+1} \dots a_{N+(M-1)}}{10^M - 1} \in \mathbb{Q}. \quad \square$$

3.4 $\mathbb{Q}$ als kleinster Körper über $\mathbb{Z}$

Wie in diesem Kapitel bewiesen wurde, haben wir durch die Erweiterung auf die rationalen Zahlen erstmals einen Körper erhalten. Auch die in den nächsten Kapiteln folgenden Erweiterungen auf die reellen und komplexen Zahlen werden auf eine Körperstruktur führen. Die rationalen Zahlen nehmen in dieser Kette von Erweiterungen insofern eine Sonderstellung ein, als da sie der *kleinste* Körper über den ganzen Zahlen sind.

Wie lässt sich „kleinster Körper“ mathematisch genau formulieren? Angenommen, die ganzen Zahlen ließen sich in einen Körper K über die Funktion $\psi : \mathbb{Z} \rightarrow K$ einbetten. Falls wir nun in K auch die rationalen Zahlen über eine Funktion $\chi : \mathbb{Q} \rightarrow K$ einbetten könnten, so können wir uns $\mathbb{Q}$ tatsächlich als kleinsten Körper über $\mathbb{Z}$ vorstellen, denn in jedem Körper K über den ganzen Zahlen sind auch die rationalen Zahlen „enthalten“. Abbildung 3.2 veranschaulicht die Situation grafisch.

In diesem Abschnitt soll mit der Funktion $\varphi : \mathbb{Z} \rightarrow \mathbb{Q}$, $\varphi(m) = \frac{m}{1}$ stets die Einbettung der ganzen Zahlen in die rationalen Zahlen gemeint sein.

Um zu beweisen, dass die rationalen Zahlen den kleinsten Körper über den ganzen Zahlen bilden, benötigen wir zunächst folgendes Lemma über bestimmte Eigenschaften einer Einbettung.

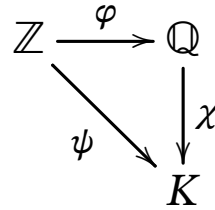


Abbildung 3.2: Lässt sich in den Körper K die Menge der ganzen Zahlen $\mathbb{Z}$ über ψ einbetten, so können wir auch $\mathbb{Q}$ in K über χ einbetten.

Lemma 3.4.1. Seien $(R_1, +, \cdot, 0_{R_1}, 1_{R_1})$ und $(R_2, \oplus, \odot, 0_{R_2}, 1_{R_2})$ zwei kommutative Halbringe mit Eins, und $\psi : R_1 \rightarrow R_2$ eine Einbettung von R_1 in R_2 . Dann gelten folgende Eigenschaften:

- (1) Es ist $\psi(0_{R_1}) = 0_{R_2}$ und $\psi(1_{R_1}) = 1_{R_2}$.
- (2) Hat $a \in R_1$ das additive Inverse $-a$, so ist $\psi(-a) = \ominus\psi(a)$.
- (3) Ist $a \in R_1$ bezüglich der Multiplikation invertierbar, so gilt das auch für $\psi(a)$ in R_2 . Es ist dann $\psi(a^{-1_{R_1}}) = \psi(a)^{\ominus 1_{R_2}}$.

Beweis. Zu (1): Für $a \in R_1$ ist $\psi(a) = \psi(a + 0_{R_1}) = \psi(a) \oplus \psi(0_{R_1})$. Da das Nullelement in einer Gruppe eindeutig ist, muss $\psi(0_{R_1}) = 0_{R_2}$ gelten. Analog lässt sich zeigen, dass $\psi(1_{R_1}) = 1_{R_2}$ ist.

Zu (2): Es ist $0_{R_2} = \psi(0_{R_1}) = \psi(a + (-a)) = \psi(a) \oplus \psi(-a)$. Da Inverse eindeutig sind, ist $\psi(-a) = \ominus\psi(a)$.

Zu (3): Angenommen, es wäre $a \cdot a^{-1_{R_1}} = 1_{R_1}$. Es ist dann $1_{R_2} = \psi(1_{R_1}) = \psi(a \cdot a^{-1_{R_1}}) = \psi(a) \odot \psi(a^{-1_{R_1}})$. Daher ist auch $\psi(a)$ in R_2 bezüglich der Multiplikation invertierbar, und es ist $\psi(a^{-1_{R_1}}) = \psi(a)^{\ominus 1_{R_2}}$. $\square$

Satz 3.4.2. Sei $(K, \oplus, \odot, 0_K, 1_K)$ ein Körper, in den $\mathbb{Z}$ über ψ eingebettet ist. Dann gibt es eine Einbettung χ von $\mathbb{Q}$ in K , sodass $\psi = \chi \circ \varphi$ ist.

Beweis. Überlegen wir zunächst salopp und etwas ungenau: Ist χ eine Einbettung von $\mathbb{Q}$ in K , und $\psi = \chi \circ \varphi$, so muss für jede Zahl $m \in \mathbb{Z}$ gelten, dass $\chi(m) = \chi(\varphi(m)) =$

$\chi\left(\frac{m}{1}\right) = \psi(m)$ ist. Zudem können wir die rationale Zahl $\frac{a}{b}$ auch als $a \cdot b^{-1}$ anschreiben. Dann ist

$$\chi\left(\frac{a}{b}\right) = \chi(a \cdot b^{-1}) = \chi(a) \odot \chi(b^{-1}) = \chi(a) \odot \chi(b)^{\ominus 1_K} = \psi(a) \odot \psi(b)^{\ominus 1_K}.$$

Es liegt daher nahe, die gesuchte Einbettung $\chi : \mathbb{Q} \rightarrow \mathbb{K}$ durch

$$\chi\left(\frac{a}{b}\right) := \psi(a) \odot \psi(b)^{\ominus 1_K}$$

zu definieren. Wir müssen nun zeigen, dass χ wohldefiniert und eine Einbettung im Sinn von Definition 2.4.1 ist.

Angenommen, $\frac{a}{b} = \frac{c}{d}$, also $a \cdot d = c \cdot b$. Es gilt damit folgende Kette von Äquivalenzen:

$$\begin{aligned} a \cdot d &= c \cdot b \Leftrightarrow \\ \psi(a \cdot d) &= \psi(c \cdot b) \Leftrightarrow \\ \psi(a) \odot \psi(d) &= \psi(c) \odot \psi(b) \Leftrightarrow \\ \psi(a) \odot \psi(b)^{\ominus 1_K} &= \psi(c) \odot \psi(d)^{\ominus 1_K} \Leftrightarrow \\ \chi\left(\frac{a}{b}\right) &= \chi\left(\frac{c}{d}\right). \end{aligned}$$

Damit ist χ wohldefiniert. Für $\frac{a}{b}, \frac{c}{d} \in \mathbb{Q}$ ist

$$\begin{aligned} \chi\left(\frac{a}{b} + \frac{c}{d}\right) &= \chi\left(\frac{a \cdot d + c \cdot b}{b \cdot d}\right) = \\ &= \psi(a \cdot d + c \cdot b) \odot \psi(b \cdot d)^{\ominus 1_K} = \\ &= (\psi(a) \odot \psi(d) \oplus \psi(c) \odot \psi(b)) \odot \psi(b)^{\ominus 1_K} \odot \psi(d)^{\ominus 1_K} = \\ &= \psi(a) \odot \psi(b)^{\ominus 1_K} \oplus \psi(c) \odot \psi(d)^{\ominus 1_K} = \\ &= \chi\left(\frac{a}{b}\right) \oplus \chi\left(\frac{c}{d}\right). \end{aligned}$$

Für die Umformungen wurde mehrmals Lemma 3.4.1 verwendet, sowie die Tatsache dass es sich bei K um einen Körper handelt und ψ eine Einbettung ist.

Weiters ist

$$\begin{aligned}
 \chi\left(\frac{a}{b} \cdot \frac{c}{d}\right) &= \chi\left(\frac{a \cdot c}{b \cdot d}\right) = \\
 &= \psi(a \cdot c) \odot \psi(b \cdot d)^{\ominus 1_K} = \\
 &= \psi(a) \odot \psi(c) \odot (\psi(b) \odot \psi(d))^{\ominus 1_K} = \\
 &= \psi(a) \odot \psi(c) \odot \psi(b)^{\ominus 1_K} \odot \psi(d)^{\ominus 1_K} = \\
 &= \psi(a) \odot \psi(b)^{\ominus 1_K} \odot \psi(c) \odot \psi(d)^{\ominus 1_K} = \\
 &= \chi\left(\frac{a}{b}\right) \odot \chi\left(\frac{c}{d}\right).
 \end{aligned}$$

Die Injektivität von χ folgt direkt aus dem Beweis der Wohldefiniertheit. Somit handelt es sich bei χ um eine Einbettung von $\mathbb{Q}$ in K . □

4 Die reellen Zahlen

Aus strukturtheoretischer Sicht sind die rationalen Zahlen als Körper ideal, und besitzen eine Vielzahl an schönen, wünschenswerten Eigenschaften. Probleme treten auf, wenn wir bestimmte nichtlineare Gleichungen wie $x^2 = 2$ lösen möchten, die in den rationalen Zahlen keine Lösung besitzen (siehe Proposition 4.2.1). Wir können außerdem feststellen: Obwohl der Zahlenstrahl eine geschlossene Linie suggeriert, enthält diese Linie zahlreiche Löcher – sogar mehr, als wir bisher zählen können!

Es ist wenig zufriedenstellend, sich mit der Existenz dieser Probleme abzufinden, denn viele Lösungen jener Probleme haben reale Entsprechungen. So tritt etwa die Zahl, deren Quadrat gleich Zwei ist, als Länge der Hypotenuse in einem gleichschenkeligen rechtwinkligen Dreieck mit Kathetenlänge Eins auf. Es wäre unbefriedigend davon zu sprechen, dass wir die Länge dieser Seite nicht beschreiben können, wenn wir durch geeignete Verfahren dazu in der Lage sind, die Länge beliebig genau zu approximieren.

Gesucht ist also eine Erweiterung der rationalen Zahlen, welche diese „Schönheitsfehler“ behebt: die reellen Zahlen. In diesem Kapitel wollen wir diese auf eine solide mathematische Grundlage stellen, wobei wir ähnlich wie in den vorigen Kapiteln vorgehen werden. Dabei wird auffallen, dass die Einführung der reellen Zahlen wesentlich mehr Aufwand bedeutet, als die der ganzen oder rationalen Zahlen. Das spiegelt sich unter anderem darin wider, dass wir in diesem Kapitel erstmals mit wichtigen Begriffen der Analysis arbeiten werden.

Die reellen Zahlen umfassen die rationalen Zahlen und ergänzen sie um die irrationalen Zahlen. Der Begriff „irrational“ stammt ursprünglich aus der Mathematik, und gelangte mit der Bedeutung „unvernünftig“ in die Alltagssprache. Das erklärt, warum wir im Titel dieser Arbeit von „unvernünftigen Zahlen“ sprechen.

4.1 Die Entdeckung der Inkommensurabilität

Für die Mathematiker der Ionischen Periode des antiken Griechenlands (600–450 v. Chr. [56, S. 150]) war die Vorstellung selbstverständlich, dass sich alle geometrischen Größen durch Zahlen oder Verhältnisse (Brüche) ausdrücken lassen [54, S. 155]. Insbesondere die Pythagoreer bauten auf dieser Vorstellung eine eigene Weltanschauung auf, nach der sich die Welt durch die Harmonie der Zahlen beschreiben lässt [56, S. 174]. Diese Philosophie spiegelt sich deutlich in dem bekannten Ausspruch „Alles ist Zahl“ wider, der von Aristoteles den Pythagoreern zugeschrieben wird [vgl. Fußnote in 53, S. 249].

Insofern entbehrt es nicht einer gewissen Ironie, dass die Entdeckung der Inkommensurabilität, also der Tatsache, dass sich *nicht* alle Strecken durch Zahlen oder Verhältnisse beschreiben lassen, ausgerechnet um 450 v. Chr. durch den Pythagoreer Hippasos von Metapont gemacht wurde [53; 56, S. 177]. Um die Geschichte dieser Entdeckung ranken sich zahlreiche Legenden. So soll Hippasos diese während einer Schifffahrt gemacht haben, und wurde von seinen erzürnten Kollegen ertränkt; eine andere Erzählung lautet, dass Hippasos nach Bekanntmachung dieser Entdeckung Schiffbruch erlitten hat, und dies als göttliche Bestrafung für die Verletzung seiner Geheimhaltungspflicht angesehen wurde [56, S. 177; 11, S. 315].

Heute geht man davon aus, dass Hippasos die Inkommensurabilität an den Seitenverhältnissen im Fünfeck entdeckt hat [11; 53]. Im Folgenden wollen wir diese Entdeckung unter Verwendung einiger elementargeometrischer Sätze skizzieren. Den Begriff „kommensurabel“ verwenden wir dabei in seiner geometrischen Deutung: Zwei Längen $\overline{AB}$ und $\overline{CD}$ heißen kommensurabel, wenn man ein $q \in \mathbb{Q}^+$ finden kann, sodass sich $\overline{AB}$ und $\overline{CD}$ als ganzzahlige Vielfache von q anschreiben lassen.

Betrachten wir dazu das regelmäßige Fünfeck in Abbildung 4.1. Wir suchen ein gemeinsames Maß der Seiten EF und FC , sprich eine positive rationale Zahl q , die sowohl $\overline{EF}$, als auch $\overline{FC}$ ganzzahlig teilt.

Zunächst können wir beobachten: Da es sich um ein regelmäßiges Fünfeck handelt, sind alle Seiten gleich lang und alle Innenwinkel gleich groß. Nach dem Seiten-Winkel-Seiten-Satz sind daher die Dreiecke ECD und ADE kongruent, was insbesondere $\angle CED = \angle ADE$ bedeutet. Es ist somit EFD ein gleichschenkeliges Dreieck, und daher $\overline{EF} = \overline{FD}$.

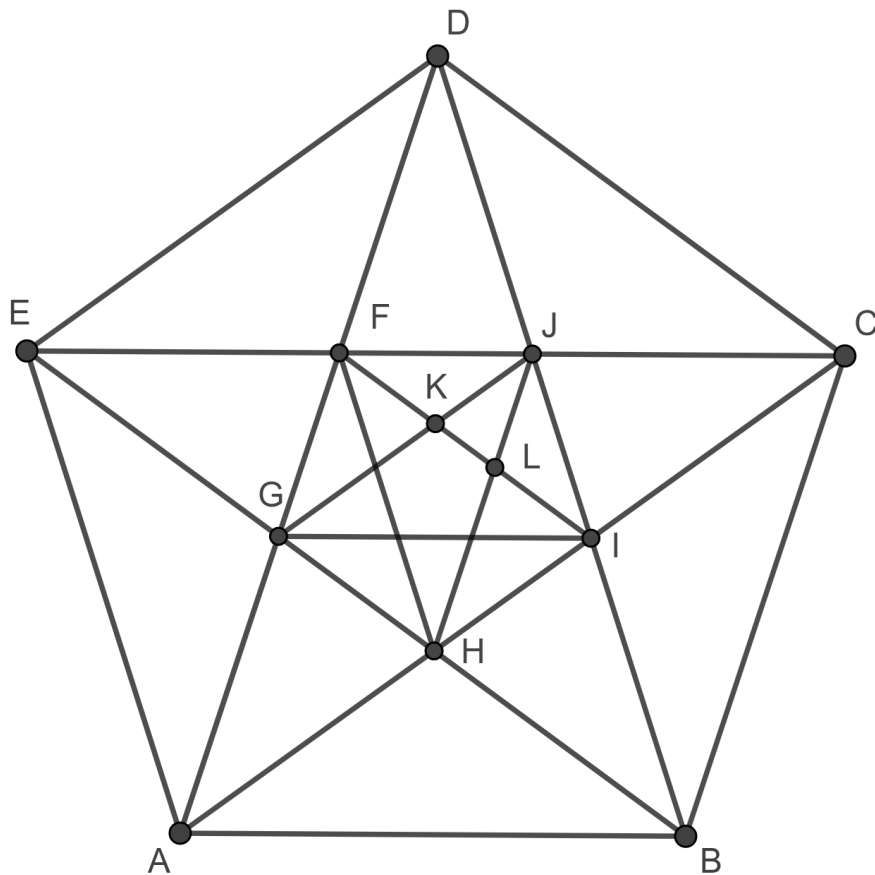


Abbildung 4.1: Ein regelmäßiges Fünfeck mit Diagonalen. Diese bilden ein inneres Fünfeck, dessen Diagonalen wieder ein Fünfeck bilden, etc.

Durch ähnliche Argumente können wir zeigen, dass die Kongruenzen $EFD \cong JCD \cong BCI \cong ABH \cong AGE$ sowie $EFG \cong DFJ \cong CJI \cong BIH \cong HAG$ gelten. Es handelt sich sicher damit auch bei dem durch die Diagonalen erzeugten inneren Fünfeck um ein regelmäßiges Fünfeck, und analog lässt sich zeigen, dass das durch die Diagonalen des inneren Fünfecks erzeugte Fünfeck wieder ein regelmäßiges Fünfeck ist. Dieser Vorgang lässt sich beliebig oft wiederholen.

Die Größe eines Innenwinkels im regelmäßigen Fünfeck beträgt 108° . Wir haben somit $108^\circ = 2 \cdot \angle EDF + \angle FDJ$, als auch $108^\circ = 180^\circ - 2 \cdot \angle DEC = 180^\circ - 2 \cdot \angle EDF$. Daraus ergibt sich $\angle EDF = 36^\circ$ und $\angle FDJ = 36^\circ$. Die Diagonalen dritteln also jeden Innenwinkel.

Damit ist $\angle DJF = 180^\circ - \angle DJC = 180^\circ - (180^\circ - 2 \cdot \angle JDC) = 2 \cdot \angle JDC = \angle EDJ$, und daher das Dreieck EDJ gleichschenkelig.

Die zuvor bewiesenen Eigenschaften gelten für alle regelmäßigen Fünfecke – insbesondere daher auch für das durch die Diagonalen erzeugte innere Fünfeck. Wir haben damit folgende Resultate hergeleitet:

- $\overline{ED} = \overline{DC} = \overline{EJ}$,
- $\overline{EF} = \overline{FD} = \overline{FI}$,
- $\overline{FJ} = \overline{JI} = \overline{KI}$.

Wären nun EF und FC kommensurabel, so gäbe es eine positive Länge q , die sowohl $\overline{EF}$ als auch $\overline{FC}$ ganzzahlig teilt. Damit teilt q auch $\overline{FC} - \overline{EF} = \overline{FJ} = \overline{KI}$. Da $\overline{EF}$ ganzzahliges Vielfaches von q ist, und zudem $\overline{EF} = \overline{FI}$ ist, gilt insgesamt, dass q auch $\overline{FI} - \overline{KI} = \overline{FK}$ teilt. Hatten wir zuvor, dass q sowohl $\overline{EF}$ wie auch $\overline{FC}$ teilt, so haben wir nun bewiesen, dass dies auch für die entsprechenden Strecken im kleineren Fünfeck gilt: q teilt $\overline{FK}$ und $\overline{KI}$ ganzzahlig.

Dieser Vorgang lässt sich nun beliebig oft wiederholen. Die Seitenlängen der inneren Fünfecke werden dabei kleiner als jede positive Zahl – und daher irgendwann auch kleiner als q . Es müsste dann q eine positive Zahl ganzzahlig teilen, die kleiner als q ist – ein Widerspruch. Somit können EF und FC nicht kommensurabel sein.

Hippasos' Entdeckung zerstörte die Auffassung, dass sich die Welt durch rationale Zahlen ausdrücken lässt. Er konnte am Fünfeck zeigen, dass es Längen gibt, die sich eben nicht durch natürliche Zahlen oder Verhältnisse beschreiben lassen – das jedoch war eine Grundidee der pythagoräischen Philosophie [53, S. 260]. Einen Ausweg aus dieser Krise fand Eudoxos von Knidos schätzungsweise um das Jahr 370 v. Chr. [53, S. 264]. Er definierte den Begriff „Verhältnis“ neu, und zwar auf eine so geschickte Weise, dass seine Definition mit allen bisher bekannten Resultaten aus der Verhältnislehre kompatibel war [56, S. 185]. Es sollte allerdings noch mehr als 2 000 Jahre dauern, bis die irrationalen Zahlen im 19. Jahrhundert auf eine solide Basis gestellt wurden. Wesentlich daran beteiligt waren unter anderem Georg Cantor und Richard Dedekind [20, S. 29].

4.2 Löcher im Zahlenstrahl

Allgemeine lineare Gleichungen der Form $a \cdot x + b = c$ mit $a \neq 0$ können wir in den rationalen Zahlen nach Proposition 3.2.6 stets lösen. An die Grenzen der rationalen Zahlen (und wie wir sehen werden sogar der reellen Zahlen) stoßen wir, wenn wir bestimmte Gleichungen höheren Grades lösen wollen.

Betrachten wir beispielsweise die Gleichung $x^2 = 2$. Nach dem pythagoräischen Lehrsatz können wir die positive Lösung dieser Gleichung als die Länge der Hypotenuse in einem gleichschenkeligen rechtwinkligen Dreieck mit Kathetenlänge Eins verstehen. In Buch 10 von Euklids Elementen wurde jedoch bewiesen, dass es kein rationales x gibt, welches diese Gleichung erfüllt [53, S. 254].

Proposition 4.2.1. *Es gibt kein $x \in \mathbb{Q}$ mit $x^2 = 2$.*

Beweis. Angenommen, es gäbe ein solches $x \in \mathbb{Q}$. Es ist dann $x \neq 0$, weswegen wir es nach Proposition 3.1.5 auf die Gestalt $\frac{m}{n}$ mit $\text{ggT}(m, n) = 1$ bringen können. Es ist somit

$$\left(\frac{m}{n}\right)^2 = \frac{m^2}{n^2} = 2.$$

Diese Gleichung können wir zu $m^2 = 2 \cdot n^2$ umformen. Nach Definition 2.5.1 ist daher Zwei ein Teiler von $m^2 = m \cdot m$. Da Zwei eine Primzahl ist, muss aufgrund von Korollar 2.5.7.1 gelten, dass Zwei bereits ein Teiler von m ist. Wir können daher $m = 2 \cdot k$ für ein geeignetes $k \in \mathbb{Z}$ schreiben.

Setzen wir das in die vorige Gleichung ein, so erhalten wir

$$m^2 = (2 \cdot k)^2 = 4 \cdot k^2 = 2 \cdot n^2.$$

Anwenden der Kürzungsregel liefert $2 \cdot k^2 = n^2$, und wie zuvor können wir daraus schließen, dass Zwei ein Teiler von n ist. Es ist somit Zwei ein Teiler von sowohl m , wie auch von n – ein Widerspruch zu $\text{ggT}(m, n) = 1$.

Es kann daher keine rationale Lösung der Gleichung $x^2 = 2$ geben. □

Bemerkenswert ist, dass wir die nicht vorhandene Lösung der Gleichung beliebig genau approximieren können. Starten wir dazu mit dem Intervall $I_0 = [1; 2]$, wobei die untere Schranke des Intervalls mit Eins gewählt wurde, da $1^2 = 1 < 2$ ist, und die obere Schranke mit Zwei, da $2^2 = 4 > 2$ ist. Wegen der Monotonie von x^2 muss die (nicht vorhandene) Lösung der Gleichung innerhalb des Intervalls I_0 liegen. Im nächsten Schritt bilden wir das arithmetische Mittel von Ober- und Untergrenze des Intervalls, und erhalten den Wert $\frac{3}{2}$. Da $\left(\frac{3}{2}\right)^2 = \frac{9}{4} > 2$ ist, definieren wir ein neues Intervall $I_1 = \left[1; \frac{3}{2}\right]$, in dem die Lösung liegen muss. Wir haben somit den Bereich, in dem die mögliche Lösung dieser Gleichung liegen kann, eingeengt.

Dieses Verfahren lässt sich wie folgt verallgemeinern:

$$I_{n+1} = \begin{cases} \left[\frac{a_n+b_n}{2}, b_n\right] & \text{wenn } \left(\frac{a_n+b_n}{2}\right)^2 \leq 2, \\ \left[a_n, \frac{a_n+b_n}{2}\right] & \text{wenn } \left(\frac{a_n+b_n}{2}\right)^2 > 2. \end{cases}$$

Wiederholtes Anwenden liefert eine Kette von Intervallen:

$$\begin{aligned} I_0 &= [1; 2] \\ I_1 &= \left[1; \frac{3}{2}\right] \\ I_2 &= \left[\frac{5}{4}; \frac{3}{2}\right] \\ I_3 &= \left[\frac{11}{8}; \frac{3}{2}\right] \\ I_4 &= \left[\frac{11}{8}; \frac{23}{16}\right] \\ &\vdots \end{aligned}$$

Während die Intervalle ineinander verschachtelt sind ($I_0 \supseteq I_1 \supseteq I_2 \supseteq \dots$), wird die Länge dieser Intervalle durch Halbierung immer kleiner; so hat etwa das Intervall I_n die Länge $\left(\frac{1}{2}\right)^n$. Wir können uns also vorstellen, dass sich die Intervalle weiter und weiter auf einen Punkt am Zahlenstrahl zusammenziehen – doch dort, wo der Punkt sein sollte, ist ein Loch.

Für die nachfolgenden Beobachtungen und den Beweis dieser Behauptung ist es notwendig, den soeben verwendeten Begriff der Intervallschachtelung zu präzisieren.

Definition 4.2.2 ([34, S. 22]). Eine Folge rationaler Intervalle $(I_n) = ([a_n, b_n])$ heißt Intervallschachtelung, falls $I_{n+1} \subseteq I_n$ für alle $n \in \mathbb{N}$ gilt, und $\lim_{n \rightarrow \infty} (b_n - a_n) = 0$ ist.

Bei dem vorigen Beispiel handelt es sich um eine Intervallschachtelung, da die Länge der Intervalle gegen Null konvergiert, und $a_n \leq \frac{a_n + b_n}{2} \leq b_n$ ist.

Lemma 4.2.3. Sei $(I_n) = ([a_n, b_n])$ eine Intervallschachtelung. Dann gilt für alle $n, m \in \mathbb{N}$: $a_n \leq b_m$. Es sind also alle Folgenglieder von (a_n) stets kleiner gleich allen Folgengliedern von (b_n) . Das lässt sich durch folgende Kette darstellen:

$$a_0 \leq a_1 \leq a_2 \leq a_3 \leq \dots \leq b_3 \leq b_2 \leq b_1 \leq b_0.$$

Beweis. Seien $n, m \in \mathbb{N}$. Wir nehmen ohne Beschränkung der Allgemeinheit an, es wäre $n \leq m$. Dann ist $I_m \subseteq I_n$, also $[a_m, b_m] \subseteq [a_n, b_n]$, woraus $a_n \leq b_m$ und $a_m \leq b_n$ folgen. Da es sich bei (I_n) um eine Intervallschachtelung handelt, ergibt sich direkt, dass (a_n) monoton steigend ist, während (b_n) monoton fällt. $\square$

Es kann höchstens eine Zahl geben, die in allen Intervallen einer Intervallschachtelung liegt.

Proposition 4.2.4. Sei $(I_n) = ([a_n, b_n])$ eine Intervallschachtelung. Dann gibt es höchstens ein $q \in \mathbb{Q}$, sodass $q \in I_n$ für alle $n \in \mathbb{N}$ ist.

Beweis. Angenommen, $p, q \in \mathbb{Q}$ liegen in allen Intervallen von (I_n) . Dann ist $0 \leq |p - q| \leq b_n - a_n$ für alle $n \in \mathbb{N}$. Wegen $\lim_{n \rightarrow \infty} (b_n - a_n) = 0$ ist $|p - q| = 0$, also $p = q$. $\square$

Proposition 4.2.5. Sei $(I_n) = ([a_n, b_n])$ eine Intervallschachtelung. Dann gibt es genau dann ein $q \in \mathbb{Q}$, sodass $q \in I_n$ für alle $n \in \mathbb{N}$ ist, wenn $q = \lim_{n \rightarrow \infty} a_n$ ist. Es ist dann ebenfalls $q = \lim_{n \rightarrow \infty} b_n$.

Beweis. Angenommen, es gibt ein $q \in \mathbb{Q}$, sodass $q \in I_n$ für alle $n \in \mathbb{N}$ ist. Sei $\varepsilon > 0$. Dann gibt es ein $N \in \mathbb{N}$, sodass $b_n - a_n < \varepsilon$ für $n \geq N$ ist. Wegen $a_n \leq q \leq b_n$ haben wir $|a_n - q| < \varepsilon$ und $|b_n - q| < \varepsilon$ für alle $n \geq N$. Es ist somit $\lim_{n \rightarrow \infty} a_n = q = \lim_{n \rightarrow \infty} b_n$.

Nehmen wir umgekehrt an, es wäre $q = \lim_{n \rightarrow \infty} a_n$. Für $\varepsilon > 0$ gibt es dann ein $N \in \mathbb{N}$, sodass $|a_n - q| < \frac{\varepsilon}{2}$ für alle $n \geq N$ ist. Wegen $\lim_{n \rightarrow \infty} (b_n - a_n) = 0$ gibt es zudem ein $M \in \mathbb{N}$, sodass für alle $m \geq M$ stets $|a_n - b_n| < \frac{\varepsilon}{2}$ ist. Sei $K = \max\{N, M\}$ und $k \geq K$. Dann ist

$$\varepsilon > |a_k - q| + |b_k - a_k| \geq |a_k - q + b_k - a_k| = |b_k - q|.$$

Somit ist $\lim_{n \rightarrow \infty} b_n = q = \lim_{n \rightarrow \infty} a_n$. Da (a_n) monoton steigend und (b_n) monoton fallend ist, muss $a_n \leq q \leq b_n$ für alle $n \in \mathbb{N}$ gelten. Es ist daher $q \in I_n$ für alle $n \in \mathbb{N}$. $\square$

Gibt es nun eine Zahl, die in allen Intervallen einer Intervallschachtelung enthalten ist, und gibt es zusätzlich eine Folge, die ab einem Index in jedem der Intervalle enthalten ist, so konvergiert die Folge gegen diese Zahl.

Proposition 4.2.6. *Sei $(I_n) = ([a_n, c_n])$ eine Intervallschachtelung und $q \in \mathbb{Q}$, sodass $q \in I_n$ für alle $n \in \mathbb{N}$ ist. Gibt es eine Folge (b_n) , sodass für alle $n \in \mathbb{N}$ ein $M \in \mathbb{N}$ existiert, mit $b_m \in I_n$ für alle $m \geq M$, so ist $\lim_{n \rightarrow \infty} b_n = q$.*

Beweis. Sei $\varepsilon > 0$. Da es sich bei (I_n) um eine Intervallschachtelung handelt, gibt es ein $N \in \mathbb{N}$, sodass $|c_N - a_N| < \varepsilon$ ist; es gilt dann $a_N \leq q \leq c_N$. Weiters gibt es ein $M(N)$, sodass $a_N \leq b_m \leq c_N$ für $m \geq M(N)$ ist. Haben wir also $n \geq \max\{N, M(N)\}$, so ist $|b_m - q| \leq |a_N - c_N| < \varepsilon$. Somit ist $\lim_{n \rightarrow \infty} b_n = q$. $\square$

Wir können nun zeigen, dass der Durchschnitt aller Intervalle des einleitenden Beispiels leer ist – oder salopp gesagt: die Intervallschachtelung zieht sich auf ein Loch am Zahlenstrahl zusammen. Denn gäbe es ein $q \in \mathbb{Q}$, das in allen Intervallen enthalten ist, so wäre nach Proposition 4.2.5 $q = \lim_{n \rightarrow \infty} a_n = \lim_{n \rightarrow \infty} b_n$. Es ist für alle $n \in \mathbb{N}$ stets $(a_n)^2 \leq 2$ und $(b_n)^2 \geq 2$, also muss das nach den Grenzwertsätzen auch für den Grenzwert q gelten: $q^2 \leq 2$ und $q^2 \geq 2$, woraus $q^2 = 2$ folgt. Eine solche rationale Zahl kann es aber nach Proposition 4.2.1 nicht geben.

Betrachten wir die Folgen (a_n) und (b_n) einer Intervallschachtelung etwas genauer. Es ist $\lim_{n \rightarrow \infty} (b_n - a_n) = 0$, was bedeutet, dass es für alle $\varepsilon > 0$ ein $N \in \mathbb{N}$ gibt, sodass für alle $n \geq N$ stets $|b_n - a_n| < \varepsilon$ gilt. Da aufgrund von Lemma 4.2.3 $a_N \leq a_n \leq b_N$ und

$a_N \leq b_n \leq b_N$ für alle $n \geq N$ gilt, muss insbesondere der Abstand aller Folgenglieder ab dem Index N kleiner als ε sein: $\forall n, m \geq N : |a_n - a_m| < \varepsilon$ und $|b_n - b_m| < \varepsilon$.

Folgen dieser Art, deren Glieder zueinander ab einem Index einen beliebig kleinen Abstand haben, werden Cauchyfolgen genannt. Sie suggerieren durch ihre Eigenschaft, dass sie stets konvergent wären – und doch haben wir soeben festgestellt, dass in den rationalen Zahlen eben nicht alle Cauchyfolgen konvergieren. Wären hingegen alle Cauchyfolgen konvergent, so hätte der Zahlenstrahl keine Löcher, denn dann zieht sich jede Intervallschachtelung auf genau einen Punkt zusammen.

Unser Ziel ist somit, die rationalen Zahlen so zu erweitern, dass jede Cauchyfolge im neuen Zahlenbereich, den reellen Zahlen, konvergent ist. Diese Eigenschaft wird auch *Vollständigkeit* genannt.

4.3 Konstruktion der reellen Zahlen

Für die Konstruktion der reellen Zahlen sind, wie im vorigen Abschnitt herausgearbeitet, Cauchyfolgen von großer Bedeutung. Im Folgenden zeigen wir daher einige wesentliche Eigenschaften solcher Folgen.

Definition 4.3.1. Sei (a_n) eine Folge in $\mathbb{Q}$. Dann heißt (a_n) Cauchyfolge, wenn

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

Die Menge aller rationalen Cauchyfolgen bezeichnen wir mit $\mathcal{C}$.

Proposition 4.3.2. Sei (a_n) eine rationale Cauchyfolge. Dann ist (a_n) beschränkt: Es gibt ein $C \in \mathbb{Q}$, sodass $|a_n| \leq C$ für alle $n \in \mathbb{N}$.

Beweis. Sei $\varepsilon = 1$. Dann gibt es ein $N \in \mathbb{N}$, sodass $|a_n - a_N| < 1$ für alle $n \geq N$ ist. Es ist dann insbesondere $|a_n| - |a_N| < 1$, also $|a_n| < 1 + |a_N|$ für alle $n \geq N$. Wir definieren $C := \max\{1 + |a_N|, |a_0|, |a_1|, \dots, |a_{N-1}|\}$. Dann ist C größer gleich allen Folgenglieder bis zum Index N , und wegen unseren vorigen Überlegungen auch größer als alle Folgenglieder ab dem Index N . Insgesamt gilt daher $|a_n| \leq C$ für alle $n \in \mathbb{N}$. $\square$

Im vorigen Abschnitt haben wir gesehen, dass jede Folge, die in einer Intervallschachtelung liegt, eine Cauchyfolge ist. Umgekehrt lässt sich zeigen, dass auch jede Cauchyfolge in einer Intervallschachtelung liegt.

Proposition 4.3.3 ([36, S. 20]). *Sei (a_n) eine rationale Cauchyfolge. Dann gibt es eine Intervallschachtelung (I_n) , sodass es für jedes $M \in \mathbb{N}$ ein $N \in \mathbb{N}$ gibt, mit $a_n \in I_M$ für alle $n \geq N$.*

Beweis. Sei $M \in \mathbb{N}$. Da es sich bei (a_n) um eine Cauchyfolge handelt, gibt es ein $K(M) \in \mathbb{N}$ mit $|a_n - a_m| < \frac{1}{M+1}$ für alle $n, m \geq K(M)$. Definieren wir daher die Intervallschachtelung (J_i) durch

$$J_i := \left[a_{K(i)} - \frac{1}{i+1}, a_{K(i)} + \frac{1}{i+1} \right],$$

so ist $a_n \in J_i$ für alle $n \geq K(i)$.

Die Intervalle aus (J_i) überlappen sich, müssen aber nicht notwendigerweise ineinander verschachtelt sein. Wir definieren daher eine neue Intervallschachtelung als den Durchschnitt $I_n := \bigcap_{i=0}^n J_i$ aller bisherigen Intervalle. Dieser ist jedenfalls ein Intervall; weiters ist offenbar $I_0 \supseteq I_1 \supseteq I_2 \supseteq \dots$. Da der Abstand der Endpunkte in I_n sicher kleiner gleich $\frac{2}{n+1}$ ist, bilden die Abstände eine Nullfolge. Es handelt sich daher bei (I_n) um eine Intervallschachtelung.

Über Induktion lässt sich nun zeigen, dass es für alle $M \in \mathbb{N}$ ein $N(M) \in \mathbb{N}$ mit $a_n \in I_M$ für alle $n \geq N(M)$ gibt. Für $M = 0$ mit $I_0 = J_0$ gilt diese Eigenschaft per Konstruktion. Gibt es nun für ein $M \in \mathbb{N}$ ein $N(M) \in \mathbb{N}$ mit $a_n \in I_M$ für alle $n \geq N(M)$, so gibt es auch einen Index $K(M+1)$, ab dem alle Folgenglieder in J_{M+1} liegen. Dann liegen aber alle Folgenglieder ab dem Index $\max\{N(M), K(M+1)\}$ in $I_M \cap J_{M+1} = I_{M+1}$.

Somit handelt es sich bei (I_n) um eine Intervallschachtelung. □

Lemma 4.3.4. *Sei (a_n) eine rationale Cauchyfolge und $q \in \mathbb{Q}$. Ist $\lim_{n \rightarrow \infty} a_n \neq q$, so gibt es ein $\varepsilon > 0$ und $N \in \mathbb{N}$ mit $|a_n - q| > \varepsilon$ für alle $n \geq N$.*

Beweis. Wegen $\lim_{n \rightarrow \infty} a_n \neq q$ gibt es ein $\eta > 0$, sodass es für alle $m \in \mathbb{N}$ ein $M(m) \geq m$ gibt mit $|a_{M(m)} - q| \geq \eta$. Da es sich bei (a_n) um eine Cauchyfolge handelt, gibt es ein $N \in \mathbb{N}$, sodass $|a_n - a_k| < \frac{\eta}{2}$ für $n, k \geq N$ ist.

Sei $n \geq N$. Es ist dann $|a_{M(n)} - q| \geq \eta$ und $M(n) \geq n \geq N$. Setzen wir $\varepsilon := \frac{\eta}{2}$ gilt

$$\eta \leq |a_{M(n)} - q| \leq \underbrace{|a_{M(n)} - a_n|}_{< \frac{\eta}{2}} + |a_n - q| < \frac{\eta}{2} + |a_n - q|,$$

und damit $|a_n - q| > \frac{\eta}{2} = \varepsilon$. □

Proposition 4.3.5. *Seien (a_n) und (b_n) rationale Cauchyfolgen. Dann sind auch die Summe $(a_n + b_n)$ und das Produkt $(a_n \cdot b_n)$ Cauchyfolgen.*

Beweis. Sei $\varepsilon > 0$. Da (a_n) und (b_n) Cauchyfolgen sind, gibt es $N, M \in \mathbb{N}$, sodass für alle $n, m \geq N$ und $k, l \geq M$ stets $|a_n - a_m| < \frac{\varepsilon}{2}$ und $|b_k - b_l| < \frac{\varepsilon}{2}$ ist. Seien $n, m \geq \max\{N, M\}$. Dann ist $\varepsilon > |a_n - a_m| + |b_n - b_m| \geq |(a_n + b_n) - (a_m + b_m)|$. Somit ist auch $(a_n + b_n)$ eine Cauchyfolge.

Zum Beweis der zweiten Eigenschaft: Nach Proposition 4.3.2 sind (a_n) und (b_n) durch eine gemeinsame Zahl C beschränkt. Sei $\varepsilon > 0$. Dann gibt es, wie wir im vorigen Schritt überlegt haben, ein $N \in \mathbb{N}$, sodass $|a_n - a_m| < \frac{\varepsilon}{2 \cdot C}$ und $|b_n - b_m| < \frac{\varepsilon}{2 \cdot C}$ für alle $n, m \geq N$ ist. Es ist dann

$$\begin{aligned} |a_n \cdot b_n - a_m \cdot b_m| &= |a_n \cdot b_n - a_m \cdot b_n + a_m \cdot b_n - a_m \cdot b_m| \leq \\ &\leq |a_n \cdot b_n - a_m \cdot b_n| + |a_m \cdot b_n - a_m \cdot b_m| = \\ &= \underbrace{|b_n|}_{\leq C} \cdot \underbrace{|a_n - a_m|}_{< \frac{\varepsilon}{2 \cdot C}} + \underbrace{|a_m|}_{\leq C} \cdot \underbrace{|b_n - b_m|}_{< \frac{\varepsilon}{2 \cdot C}} < \\ &< \varepsilon. \end{aligned}$$

Somit handelt es sich bei $(a_n \cdot b_n)$ um eine Cauchyfolge. □

Für die Konstruktion der ganzen und rationalen Zahlen haben wir stets bestimmte Elemente, die sich eine gemeinsame Eigenschaft teilen, als äquivalent angesehen. Der Weg zur Definition der reellen Zahlen führt ebenfalls über eine solche Äquivalenzrelation auf der Menge aller rationalen Cauchyfolgen. Wir wollen im Prinzip jene Cauchyfolgen zusammenfassen, die den gleichen Grenzwert besitzen. Da allerdings nicht jede Cauchyfolge in den rationalen Zahlen konvergiert, können wir nicht definieren, dass wir zwei

Cauchyfolgen als äquivalent ansehen, falls die womöglich nicht existenten Grenzwerte gleich sind. Wenn wir das vorige Beispiel zu den Intervallschachtelungen betrachten, so sehen wir, dass die Cauchyfolgen zwar nicht konvergieren, aber ihr Abstand beliebig klein wird. Diese Eigenschaft verwenden wir nun zur Definition der Äquivalenzrelation.

Definition 4.3.6. Seien (a_n) und (b_n) rationale Cauchyfolgen. Wir definieren die Relation $\sim$ auf $\mathcal{C}$ als

$$(a_n) \sim (b_n) :\Leftrightarrow \lim_{n \rightarrow \infty} (a_n - b_n) = 0.$$

Proposition 4.3.7. Bei der Relation $\sim$ auf $\mathcal{C}$ handelt es sich um eine Äquivalenzrelation. Es sind daher folgende Eigenschaften erfüllt:

- (1) $\forall (a_n) \in \mathcal{C} : (a_n) \sim (a_n)$ (Reflexivität).
- (2) $\forall (a_n), (b_n) \in \mathcal{C} : (a_n) \sim (b_n) \Rightarrow (b_n) \sim (a_n)$ (Symmetrie).
- (3) $\forall (a_n), (b_n), (c_n) \in \mathcal{C} : ((a_n) \sim (b_n) \wedge (b_n) \sim (c_n)) \Rightarrow (a_n) \sim (c_n)$ (Transitivität).

Beweis. Zur Reflexivität: Es ist $\lim_{n \rightarrow \infty} (a_n - a_n) = \lim_{n \rightarrow \infty} 0 = 0$, und damit $(a_n) \sim (a_n)$.

Zur Symmetrie: Angenommen, es wäre $(a_n) \sim (b_n)$, also $\lim_{n \rightarrow \infty} (a_n - b_n) = 0$. Wegen der Rechenregeln für Grenzwerte können wir die vorige Gleichung mit -1 multiplizieren, und erhalten $\lim_{n \rightarrow \infty} (b_n - a_n) = 0$, womit $(b_n) \sim (a_n)$ gezeigt ist.

Die Transitivität folgt ebenfalls aus den Rechenregeln für Grenzwerte. Es ist nach Voraussetzung $\lim_{n \rightarrow \infty} (a_n - b_n) = 0 = \lim_{n \rightarrow \infty} (b_n - c_n)$. Addieren dieser Ausdrücke liefert $\lim_{n \rightarrow \infty} (a_n - c_n) = 0$, also $(a_n) \sim (c_n)$. $\square$

Definition 4.3.8. Die Menge der reellen Zahlen $\mathbb{R}$ ist definiert als

$$\mathbb{R} := \mathcal{C} / \sim.$$

Die Äquivalenzklasse der Folge (a_n) schreiben wir als $C(a_n)$ an. Dort, wo wir hervorheben wollen, dass es sich um eine Äquivalenzklasse von Folgen handelt, schreiben wir $C((a_n))$.

4.4 Rechenoperationen auf $\mathbb{R}$

Wie in den letzten Kapiteln wollen wir die Addition und die Multiplikation auf die reellen Zahlen übertragen.

Da es sich bei den reellen Zahlen um Äquivalenzklassen von Folgen handelt, und wir Folgen addieren und multiplizieren können, bietet es sich an, für die Addition und Multiplikation auf den reellen Zahlen darauf zurückzugreifen. Proposition 4.3.5 garantiert uns, dass wir auf diese Weise wieder Cauchyfolgen erhalten.

Definition 4.4.1. Die Addition $+$: $\mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$ ist definiert durch

$$C((a_n)) + C((b_n)) := C((a_n + b_n)).$$

Da wir in der Definition Repräsentanten von Äquivalenzklassen verwendet haben, müssen wir zeigen, dass die Addition auf den reellen Zahlen wohldefiniert ist.

Proposition 4.4.2. Die Addition auf den reellen Zahlen ist wohldefiniert. Sind also $(a'_n) \sim (a_n)$ und $(b'_n) \sim (b_n)$, so ist $C(a_n) + C(b_n) = C(a'_n) + C(b'_n)$.

Beweis. Es ist nach Voraussetzung $\lim_{n \rightarrow \infty} (a'_n - a_n) = 0 = \lim_{n \rightarrow \infty} (b'_n - b_n)$. Wegen der Rechenregeln für Grenzwerte ist damit $\lim_{n \rightarrow \infty} ((a'_n + b'_n) - (a_n + b_n)) = 0$, und daher $C(a_n + b_n) = C(a'_n + b'_n)$. $\square$

Definition 4.4.3. Die Multiplikation $\cdot$: $\mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$ ist definiert durch

$$C((a_n)) \cdot C((b_n)) := C((a_n \cdot b_n)).$$

Proposition 4.4.4. Die Multiplikation auf den reellen Zahlen ist wohldefiniert. Sind also $(a'_n) \sim (a_n)$ und $(b'_n) \sim (b_n)$, so ist $C(a_n) \cdot C(b_n) = C(a'_n) \cdot C(b'_n)$.

Beweis. Bei (a'_n) und (b_n) handelt es sich um Cauchyfolgen, weswegen es nach Proposition 4.3.2 ein $C \in \mathbb{Q}$ gibt, sodass $|a'_n| \leq C$ und $|b_n| \leq C$ für alle $n \in \mathbb{N}$ gilt.

Sei $\varepsilon > 0$. Da $\lim_{n \rightarrow \infty} (a'_n - a_n) = 0 = \lim_{n \rightarrow \infty} (b'_n - b_n)$ ist, gibt es ein $N \in \mathbb{N}$, sodass für alle $n \geq N$ stets $|a'_n - a_n| < \frac{\varepsilon}{2 \cdot C}$ und $|b'_n - b_n| < \frac{\varepsilon}{2 \cdot C}$ erfüllt ist. Es gilt damit:

$$\begin{aligned} |a'_n \cdot b'_n - a_n \cdot b_n| &= |a'_n \cdot b'_n - a'_n \cdot b_n + a'_n \cdot b_n - a_n \cdot b_n| \leq \\ &\leq |a'_n \cdot b'_n - a'_n \cdot b_n| + |a'_n \cdot b_n - a_n \cdot b_n| = \\ &= \underbrace{|a'_n|}_{\leq C} \cdot \underbrace{|b'_n - b_n|}_{< \frac{\varepsilon}{2 \cdot C}} + \underbrace{|b_n|}_{\leq C} \cdot \underbrace{|a'_n - a_n|}_{< \frac{\varepsilon}{2 \cdot C}} < \\ &< \varepsilon. \end{aligned}$$

Somit ist $\lim_{n \rightarrow \infty} (a'_n \cdot b'_n - a_n \cdot b_n) = 0$, und $C(a_n \cdot b_n) = C(a'_n \cdot b'_n)$. □

Die reellen Zahlen bilden einen Körper. Als Null- und Einselement verwenden wir die Äquivalenzklassen der konstanten Folgen $(0)_{n \in \mathbb{N}}$ und $(1)_{n \in \mathbb{N}}$.

Proposition 4.4.5. $(\mathbb{R}, +, \cdot, C(0), C(1))$ ist ein Körper. Es sind daher die Körperaxiome erfüllt (vgl. [50, S. 264]):

- (1) $\forall a, b, c \in \mathbb{R} : (a + b) + c = a + (b + c)$ (Assoziativität von $+$).
- (2) $\forall a, b \in \mathbb{R} : a + b = b + a$ (Kommutativität von $+$).
- (3) $\forall a \in \mathbb{R} : a + C(0) = a$ (Nullelement).
- (4) $\forall a \in \mathbb{R} : \exists -a \in \mathbb{R} : a + (-a) = C(0)$ (Invertierbarkeit von $+$).
- (5) $\forall a, b, c \in \mathbb{R} : (a \cdot b) \cdot c = a \cdot (b \cdot c)$ (Assoziativität von $\cdot$).
- (6) $\forall a, b \in \mathbb{R} : a \cdot b = b \cdot a$ (Kommutativität von $\cdot$).
- (7) $\forall a \in \mathbb{R} : a \cdot C(1) = a$ (Einselement).
- (8) $\forall a \in \mathbb{R} \setminus \{C(0)\} : \exists a^{-1} \in \mathbb{R} : a \cdot a^{-1} = C(1)$ (Invertierbarkeit von $\cdot$).
- (9) $\forall a, b, c \in \mathbb{R} : a \cdot (b + c) = a \cdot b + a \cdot c$ (Distributivität).

Beweis. Seien $a = C(a_n)$, $b = C(b_n)$ und $c = C(c_n)$ reelle Zahlen. Für den Beweis der Assoziativität der Addition greifen wir auf die Assoziativität der rationalen Zahlen zurück:

$$\begin{aligned}
 (a + b) + c &= (C(a_n) + C(b_n)) + C(c_n) = \\
 &= C(a_n + b_n) + C(c_n) = \\
 &= C((a_n + b_n) + c_n) = \\
 &= C(a_n + (b_n + c_n)) = \\
 &= C(a_n) + C(b_n + c_n) = \\
 &= C(a_n) + (C(b_n) + C(c_n)) = \\
 &= a + (b + c).
 \end{aligned}$$

Auch die Kommutativität der Addition lässt sich auf die rationalen Zahlen zurückführen:

$$a + b = C(a_n) + C(b_n) = C(a_n + b_n) = C(b_n + a_n) = C(b_n) + C(a_n) = b + a.$$

Analog die Invertierbarkeit der Addition, wobei wir $-a := C(-a_n)$ setzen:

$$a + (-a) = C(a_n) + C(-a_n) = C(a_n + (-a_n)) = C(0).$$

Für die Assoziativität der Multiplikation greifen wir wieder auf die rationalen Zahlen zurück. Es ist

$$\begin{aligned}
 (a \cdot b) \cdot c &= (C(a_n) \cdot C(b_n)) \cdot C(c_n) = \\
 &= C(a_n \cdot b_n) \cdot C(c_n) = \\
 &= C((a_n \cdot b_n) \cdot c_n) = \\
 &= C(a_n \cdot (b_n \cdot c_n)) = \\
 &= C(a_n) \cdot C(b_n \cdot c_n) = \\
 &= C(a_n) \cdot (C(b_n) \cdot C(c_n)) = \\
 &= a \cdot (b \cdot c).
 \end{aligned}$$

Zur Kommutativität der Multiplikation:

$$a \cdot b = C(a_n) \cdot C(b_n) = C(a_n \cdot b_n) = C(b_n \cdot a_n) = C(b_n) \cdot C(a_n) = b \cdot a.$$

Für den Beweis der Invertierbarkeit der Multiplikation ist Vorsicht nötig. Wir können nicht einfach den Kehrwert der Folge verwenden, da einzelne Folgenglieder gleich Null sein können, selbst wenn die Folge nicht gegen Null konvergiert. Da die Folge aber keine Nullfolge ist, muss es nach Lemma 4.3.4 einen Index $N \in \mathbb{N}$ geben, ab dem die Folgenglieder von Null verschieden sind. Wir definieren daher die Folge (b_n) durch

$$b_n = \begin{cases} a_n & \text{wenn } n < N, \\ \frac{1}{a_n} & \text{sonst.} \end{cases}$$

Sei $\varepsilon > 0$. Es ist dann für alle $n \geq N$ stets $|a_n \cdot b_n - 1| = 0 < \varepsilon$, womit $C(a_n) \cdot C(b_n) = C(1)$ gezeigt ist.

Zuletzt die Distributivität, die wir wieder auf die rationalen Zahlen zurückführen können:

$$\begin{aligned} a \cdot (b + c) &= C(a_n) \cdot (C(b_n) + C(c_n)) = \\ &= C(a_n) \cdot C(b_n + c_n) = \\ &= C(a_n \cdot (b_n + c_n)) = \\ &= C(a_n \cdot b_n + a_n \cdot c_n) = \\ &= C(a_n \cdot b_n) + C(a_n \cdot c_n) = \\ &= C(a_n) \cdot C(b_n) + C(a_n) \cdot C(c_n) = \\ &= a \cdot b + a \cdot c. \end{aligned} \quad \square$$

Das additive Inverse von a schreiben wir wie üblich als $-a$ an, während wir das multiplikative Inverse von $a \neq C(0)$ als a^{-1} bezeichnen. Zudem vereinbaren wir, anstelle von $a + (-b)$ direkt $a - b$ zu schreiben.

Da die reellen Zahlen einen Körper bilden, folgen insbesondere die Eigenschaften aus Proposition 3.2.6.

Um eine Ordnung auf den reellen Zahlen zu definieren, gehen wir ähnlich vor wie in den rationalen Zahlen: Eine reelle Zahl a soll kleiner gleich einer Zahl b sein, wenn es eine positive reelle Zahl c gibt, sodass $a + c = b$ ist. Dazu benötigen wir zunächst eine Definition von „positiv“ in den reellen Zahlen.

Wir haben die reellen Zahlen über Äquivalenzklassen von Cauchyfolgen definiert. Wenn nun die Folge (a_n) der reellen Zahl $C(a_n)$ ab einem gewissen Index stets größer als eine positive rationale Zahl ist, so wollen wir auch die zugehörige reelle Zahl als positiv bezeichnen. Die Idee dahinter ist, dass der Grenzwert dieser Folge, oder das Loch im Zahlenstrahl, dann ebenfalls auf der positiven Seite des Zahlenstrahls liegen muss.

Die Definition „größer als eine positive rationale Zahl“ ist notwendig – es reicht nicht aus zu fordern, dass die Folge ab einem Index größer als Null ist. Das wird am Beispiel $C(\frac{1}{n}) = C(0)$ deutlich: Obwohl $\frac{1}{n}$ für alle $n \in \mathbb{N}$ stets größer als Null ist, entspricht die zugehörige reelle Zahl dem Nullelement, welches wir als weder positiv noch negativ ansehen wollen.

Definition 4.4.6. Wir nennen die reelle Zahl $C(a_n)$ positiv, falls es ein $c \in \mathbb{Q}^+$ und ein $N \in \mathbb{N}$ gibt, sodass $a_n \geq c$ für alle $n \geq N$ ist. Ist $a_n \leq c$ für ein $c \in \mathbb{Q}^-$ und alle $n \geq N$, so nennen wir $C(a_n)$ negativ. Das Nullelement $C(0)$ wird als weder positiv noch negativ definiert.

Zudem definieren wir

$$\begin{aligned}\mathbb{R}^+ &:= \{C(a_n) \in \mathbb{R} \mid \exists c \in \mathbb{Q}^+ : \exists N \in \mathbb{N} : \forall n \geq N : a_n \geq c\}, \\ \mathbb{R}^- &:= \{C(a_n) \in \mathbb{R} \mid \exists c \in \mathbb{Q}^- : \exists N \in \mathbb{N} : \forall n \geq N : a_n \leq c\}\end{aligned}$$

als die Menge der positiven, beziehungsweise der negativen reellen Zahlen.

Da diese Definition Repräsentanten von Äquivalenzklassen verwendet, müssen wir die Wohldefiniertheit beweisen.

Proposition 4.4.7. Die Eigenschaften „positiv“ und „negativ“ von reellen Zahlen sind wohldefiniert. Ist also $C(a_n)$ positiv und $(b_n) \sim (a_n)$, so ist auch $C(b_n)$ positiv. Analoges gilt für negative reelle Zahlen.

Beweis. Angenommen, $C(a_n)$ ist positiv und $(b_n) \sim (a_n)$. Es gibt somit ein $c \in \mathbb{Q}^+$ und ein $N \in \mathbb{N}$, sodass $a_n \geq c$ für alle $n \geq N$ ist. Da $\lim_{n \rightarrow \infty} (a_n - b_n) = 0$ ist, gibt es ein $M \in \mathbb{N}$, sodass $|a_m - b_m| < \frac{c}{2}$ für alle $m \geq M$ ist. Sei nun $K = \max\{N, M\}$ und $k \geq K$. Es ist dann

$$c - b_k \leq a_k - b_k \leq |a_k - b_k| < \frac{c}{2},$$

und daher $0 < \frac{c}{2} < b_k$ für alle $k \geq K$. Somit ist auch $C(b_n)$ positiv.

Der Beweis für die Wohldefiniertheit der Negativität läuft analog: Ist $C(a_n)$ negativ und $(b_n) \sim (a_n)$, so gibt es ein $c \in \mathbb{Q}^-$ und ein $N \in \mathbb{N}$, sodass $a_n \leq c$ für alle $n \geq N$. Wegen $\lim_{n \rightarrow \infty} (a_n - b_n) = 0$ gibt es ein $M \in \mathbb{N}$, sodass $|b_m - a_m| < -\frac{c}{2}$ für alle $m \geq M$ ist. Setzen wir $K = \max\{N, M\}$, und ist $k \geq K$, so haben wir

$$b_k - c \leq b_k - a_k \leq |b_k - a_k| < -\frac{c}{2},$$

und damit $b_k < \frac{c}{2} < 0$. Somit ist $C(b_n)$ negativ. □

Proposition 4.4.8. *Es ist $\mathbb{R} = \mathbb{R}^+ \cup \{C(0)\} \cup \mathbb{R}^-$ und $\mathbb{R}^+ \cap \mathbb{R}^- = \emptyset$. Somit ist jede reelle Zahl entweder positiv, negativ, oder gleich dem Nullelement.*

Beweis. Offensichtlich ist $C(0)$ Element der Vereinigung. Für $C(a_n) \in \mathbb{R}$ mit $C(a_n) \neq C(0)$ ist (a_n) nicht konvergent, oder $\lim_{n \rightarrow \infty} a_n \neq 0$. Nach Lemma 4.3.4 gibt es daher ein $\varepsilon > 0$ und $N \in \mathbb{N}$ mit $|a_n| > \varepsilon$ für alle $n \geq N$. Da (a_n) eine Cauchyfolge ist, gibt es weiters einen Index $M \in \mathbb{N}$, ab dem alle Folgenglieder entweder positiv oder negativ sind. Es ist dann also $a_n < -\varepsilon$ oder $a_n > \varepsilon$ für $n \geq \max\{N, M\}$, und somit $C(a_n)$ entsprechend positiv oder negativ.

Nehmen wir nun an, es gäbe eine reelle Zahl $C(a_n)$ die sowohl positiv als auch negativ ist. Es gibt daher ein $c \in \mathbb{Q}^+$, ein $d \in \mathbb{Q}^-$, sowie $N, M \in \mathbb{N}$, sodass für alle $n \geq N$ und $m \geq M$ stets $a_n \geq c$ und $a_m \leq d$ gilt. Sei nun $n \geq \max\{M, N\}$. Dann wäre $a_n \leq d < 0 < c \leq a_n$, ein Widerspruch. Somit erhalten wir insgesamt, dass jede reelle Zahl entweder positiv, negativ, oder gleich dem Nullelement ist. □

Lemma 4.4.9. *Sei $C(a_n) \in \mathbb{R}$. Falls $C(a_n)$ positiv ist, so ist $-C(a_n)$ negativ. Ist $C(a_n)$ negativ, so ist $-C(a_n)$ positiv.*

Beweis. Angenommen $C(a_n)$ ist positiv, dann gibt es ein $c \in \mathbb{Q}^+$ und ein $N \in \mathbb{N}$ mit $a_n \geq c$ für alle $n \geq N$. Multiplizieren dieser Gleichung mit -1 liefert $-a_n \leq -c$ für alle $n \geq N$. Da $-c$ nach Lemma 3.2.10 negativ ist, ist auch $C(-a_n) = -C(a_n)$ negativ.

Die zweite Aussage folgt analog. $\square$

Auch in den reellen Zahlen gelten die uns bekannten Vorzeichenregeln (siehe Seite 49). Ist beispielsweise $C(a_n) \in \mathbb{R}^+$ und $C(b_n) \in \mathbb{R}^-$, so gibt es ein $c \in \mathbb{Q}^+$, ein $d \in \mathbb{Q}^-$ sowie einen gemeinsamen Index $N \in \mathbb{N}$, sodass $a_n \geq c$ und $b_n \leq d$ für alle $n \geq N$ ist. Es ist dann $a_n \cdot b_n \leq a_n \cdot d \leq c \cdot d$, wobei $c \cdot d$ nach den Rechenregeln für rationale Zahlen negativ ist. Somit ist $C(a_n) \cdot C(b_n)$ eine negative reelle Zahl.

Definition 4.4.10. Seien $a, b \in \mathbb{R}$. Dann definieren wir

$$a \leq b :\Leftrightarrow \exists c \in \mathbb{R} \setminus \mathbb{R}^- : a + c = b.$$

Zudem schreiben wir $b \geq a$ für $a \leq b$. Ist das c in obiger Definition positiv (also ungleich $C(0)$), so verwenden wir $<$ und $>$ anstelle von $\leq$ und $\geq$.

Wir erhalten abermals die Tatsache, dass die negativen reellen Zahlen kleiner als $C(0)$, und die positiven reellen Zahlen größer als $C(0)$ sind. Das Argument hierfür lässt sich exakt wie auf Seite 72 führen.

Proposition 4.4.11. Die Relation $\leq$ auf $\mathbb{R}$ ist eine Totalordnung. Für alle $a, b \in \mathbb{R}$ gilt daher $a \leq b$ oder $b \leq a$. Zudem sind die Eigenschaften einer Ordnungsrelation erfüllt:

- (1) $\forall a \in \mathbb{R} : a \leq a$ (Reflexivität).
- (2) $\forall a, b \in \mathbb{R} : (a \leq b \wedge b \leq a) \Rightarrow a = b$ (Antisymmetrie).
- (3) $\forall a, b, c \in \mathbb{R} : (a \leq b \wedge b \leq c) \Rightarrow a \leq c$ (Transitivität).

Beweis. Die Reflexivität folgt direkt aus $a + C(0) = a$. Für die Antisymmetrie nehmen wir an, dass $a \leq b$ und $b \leq a$ wäre. Es gibt daher $c, c' \in \mathbb{R} \setminus \mathbb{R}^-$, sodass $b = a + c$ und $a = b + c'$ ist. Einsetzen der ersten Gleichung in die zweite liefert $a = a + (c + c')$, also $c + c' = C(0)$. Wären c oder c' positiv, so müsste es auch die Summe $c + c' = C(0)$ sein, ein Widerspruch. Somit ist $c = c' = C(0)$ und daher $a = b$.

Zur Transitivität: Ist $a \leq b$ und $b \leq c$, so gibt es $d, d' \in \mathbb{R} \setminus \mathbb{R}^-$, sodass $b = a + d$ und $c = b + d'$ ist. Es ist dann $c = a + (d + d')$, wobei $d + d'$ nach unseren Überlegungen eine Summe nichtnegativer reeller Zahlen ist, also selbst wieder nichtnegativ ist. Somit ist $a \leq c$.

Um zu zeigen, dass es sich um eine Totalordnung handelt, betrachten wir $a, b \in \mathbb{R}$. Nach Proposition 3.2.6 hat die Gleichung $a + x = b$ stets eine eindeutige Lösung $x \in \mathbb{R}$, die nach Proposition 4.4.8 entweder negativ oder nichtnegativ ist. Falls die Lösung nichtnegativ ist, so folgt direkt aus Definition 4.4.10 $a \leq b$. Ist x hingegen negativ, so können wir auf beiden Seiten der Gleichung $-x \in \mathbb{R}^+$ addieren, und erhalten $a = b + (-x)$, also $b \leq a$. Somit sind stets zwei reelle Zahlen über $\leq$ vergleichbar. $\square$

Proposition 4.4.12. *Die Relation $\leq$ auf den reellen Zahlen erfüllt die Ordnungsaxiome:*

- (1) Für $a, b, c \in \mathbb{R}$ ist $a \leq b \Leftrightarrow a + c \leq b + c$.
- (2) Für $a, b \in \mathbb{R}$ mit $a > C(0)$ und $b > C(0)$ ist $a \cdot b > C(0)$.

Beweis. Es ist $a \leq b$ äquivalent dazu, dass es ein $d \in \mathbb{R} \setminus \mathbb{R}^-$ mit $b = a + d$ gibt. Diese Gleichung ist wiederum äquivalent zu $b + c = (a + c) + d$, also $a + c \leq b + c$.

Für den Beweis der zweiten Behauptung nehmen wir an, es wäre $a > C(0)$ und $b > C(0)$. Nach unseren vorigen Beobachtungen ist das Produkt zweier positiver reeller Zahlen wieder positiv, und da alle positiven reellen Zahlen größer als $C(0)$ sind, folgt $a \cdot b > C(0)$. $\square$

Wie in den rationalen Zahlen dürfen wir auch in den reellen Zahlen mit Ungleichungen rechnen.

Proposition 4.4.13. *Seien $a, b, c, d \in \mathbb{R}$. Dann gilt:*

- (1) Ist $a \leq b$ und $c \leq d$, so ist $a + c \leq b + d$.
- (2) Ist c positiv, so ist $a \leq b$ äquivalent zu $a \cdot c \leq b \cdot c$.
- (3) Ist c negativ, so ist $a \leq b$ äquivalent zu $b \cdot c \leq a \cdot c$.

Beweis. Zu (1): Ist $a \leq b$ und $c \leq d$, so gibt es $r, s \in \mathbb{R} \setminus \mathbb{R}^-$ mit $b = a + r$ und $d = c + s$. Addieren dieser Gleichungen liefert $b + d = a + c + (r + s)$. Da die Summe zweier nichtnegativer reeller Zahlen wieder nichtnegativ ist, folgt daraus $a + c \leq b + d$.

Zu (2): Angenommen, c wäre positiv und $a \leq b$. Dann gibt es ein $r \in \mathbb{R}$ mit $b = a + r$. Multiplizieren dieser Gleichung mit c liefert $b \cdot c = a \cdot c + (r \cdot c)$. Da es sich bei r und c um nichtnegative reelle Zahlen handelt, ist nach unseren vorigen Beobachtungen auch ihr Produkt nichtnegativ, und daher $a \cdot c \leq b \cdot c$.

Ist umgekehrt $a \cdot c \leq b \cdot c$, so gibt es ein $r \in \mathbb{R} \setminus \mathbb{R}^-$ mit $b \cdot c = a \cdot c + r$. Multiplizieren wir diese Gleichung mit c^{-1} , so erhalten wir $b = a + r \cdot c^{-1}$. Da c positiv ist, muss auch c^{-1} positiv sein, und damit auch das Produkt $r \cdot c^{-1}$. Es ist daher $a \leq b$.

Zu (3): Ist c negativ, so muss $-c$ nach Lemma 4.4.9 positiv sein. Wegen Punkt (2) ist $a \leq b$ äquivalent zu $a \cdot (-c) \leq b \cdot (-c)$. Es gibt daher ein $r \in \mathbb{R} \setminus \mathbb{R}^-$ mit $b \cdot (-c) = a \cdot (-c) + r$. Multiplizieren mit $C(-1)$ liefert $b \cdot c = a \cdot c - r$, was zu $a \cdot c = b \cdot c + r$ äquivalent ist, also $b \cdot c \leq a \cdot c$. $\square$

Proposition 4.4.14. Seien $C(a_n)$ und $C(b_n)$ reelle Zahlen. Dann folgt aus $C(a_n) < C(b_n)$, dass es ein $N \in \mathbb{N}$ mit $a_n < b_n$ für alle $n \geq N$ gibt. Sind umgekehrt (a_n) und (b_n) rationale Cauchyfolgen, und gibt es ein $N \in \mathbb{N}$ mit $a_n < b_n$ für alle $n \geq N$, so ist $C(a_n) \leq C(b_n)$.

Beweis. Angenommen, es wäre $C(a_n) < C(b_n)$. Dann ist $C(b_n - a_n)$ positiv, weswegen es ein $c \in \mathbb{Q}^+$ und ein $N \in \mathbb{N}$ mit $b_n - a_n \geq c$ für alle $n \geq N$ gibt. Es ist daher $b_n \geq c + a_n > a_n$ für alle $n \geq N$.

Für den Beweis der zweiten Aussage betrachten wir die Folge

$$c_n := \begin{cases} 0 & \text{wenn } n < N, \\ b_n - a_n & \text{sonst.} \end{cases}$$

Die Folge (c_n) ist eine rationale Cauchyfolge, deren Glieder nichtnegativ sind. Daher ist auch $C(c_n)$ nichtnegativ, denn sonst müsste (c_n) ab einem Index kleiner als eine feste negative Zahl sein, im Widerspruch zur Nichtnegativität. Damit folgt, wegen $C(a_n) + C(c_n) = C(b_n)$, dass $C(a_n) \leq C(b_n)$ ist. $\square$

Wie wir es aus den letzten Erweiterungen gewohnt sind, stellt auch die Ordnung in den reellen Zahlen keine Wohlordnung dar. Das wird am Beispiel der Menge der negativen reellen Zahlen $\mathbb{R}^-$ deutlich, die nichtleer ist, allerdings kein kleinstes Element besitzt.

4.5 Vollständigkeit von $\mathbb{R}$

In Abschnitt 4.2 haben wir gezeigt, dass der Zahlenstrahl der rationalen Zahlen unerwartete Löcher enthält, und festgestellt, dass die Existenz dieser Löcher an die Konvergenz von Cauchyfolgen geknüpft ist. Konvergiert in einem metrischen Raum jede Cauchyfolge, so nennt man diesen Raum vollständig. Wir wollen nun zeigen, dass die reellen Zahlen, wie wir sie definiert haben, tatsächlich vollständig sind.

Zunächst benötigen wir den Begriff der rationalen reellen Zahl. Wie der Name bereits andeutet, handelt es sich bei diesen um die in die reellen Zahlen eingebetteten rationalen Zahlen.

Definition 4.5.1. Wir nennen $\alpha \in \mathbb{R}$ eine rationale reelle Zahl, wenn es ein $a \in \mathbb{Q}$ mit $\alpha = C((a)_{n \in \mathbb{N}})$ gibt. Existiert es ein solches a nicht, so nennen wir α irrational.

Beispiele für rationale reelle Zahlen sind das Nullelement $C(0)$, oder das Einselement $C(1)$ der reellen Zahlen. Eine irrationale reelle Zahl ist etwa jene reelle Zahl, die durch die Cauchyfolgen jener Intervallschachtelung definiert wird, die wir in Abschnitt 4.2 untersucht haben.

Für die reellen Zahlen definieren wir die Betragsfunktion wie in den rationalen Zahlen.

Definition 4.5.2. Für $\alpha \in \mathbb{R}$ definieren wir

$$|\alpha| := \begin{cases} \alpha & \text{wenn } \alpha \geq C(0), \\ -\alpha & \text{sonst.} \end{cases}$$

Die Betragsfunktion weist alle bekannten Eigenschaften, unter anderem die Produktregel und die Dreiecksungleichung, auf. Aus der Definition geht mit Hilfe von Lemma 4.4.9 hervor, dass der Betrag stets eine nichtnegative reelle Zahl ist.

Proposition 4.5.3. Sei $\alpha = C((a_n)_{n \in \mathbb{N}})$ eine reelle Zahl. Dann ist

$$|C((a_n)_{n \in \mathbb{N}})| = C(|a_n|_{n \in \mathbb{N}}).$$

Beweis. Behandeln wir zunächst den Fall $C((a_n)_{n \in \mathbb{N}}) = C(0)$. Dann ist $|C((a_n)_{n \in \mathbb{N}})| = C(0)$, und nach Definition 4.3.6 auch $\lim_{n \rightarrow \infty} a_n = 0$, woraus $\lim_{n \rightarrow \infty} |a_n| = 0$ folgt. Somit ist $C((a_n)_{n \in \mathbb{N}}) = C(|a_n|_{n \in \mathbb{N}})$.

Ist $C((a_n)_{n \in \mathbb{N}})$ positiv, so gibt es nach Definition 4.4.6 ein $c \in \mathbb{Q}^+$ und einen Index N , sodass $a_n \geq c$ für alle $n \geq N$ erfüllt ist. Ab diesem Index ist $a_n - |a_n| = 0$, und damit $|C((a_n)_{n \in \mathbb{N}})| = C(|a_n|_{n \in \mathbb{N}})$.

Ist $C((a_n)_{n \in \mathbb{N}})$ negativ, so gehen wir ähnlich vor: Es ist $|C((a_n)_{n \in \mathbb{N}})| = C((-a_n)_{n \in \mathbb{N}})$, und es gibt ein $c \in \mathbb{Q}^-$ sowie ein $N \in \mathbb{N}$, sodass $a_n \leq c$ für alle $n \geq N$ ist. Ab diesem Index ist daher $-a_n - |a_n| = -a_n - (-a_n) = 0$, und somit ebenfalls $|C((a_n)_{n \in \mathbb{N}})| = C(|a_n|_{n \in \mathbb{N}})$. $\square$

Bemerkung. Für einige Propositionen und Beweise in diesem Abschnitt ist es hilfreich, sich $C(a_n) = \lim_{n \rightarrow \infty} a_n$ vorzustellen. Auf die obige Proposition angewandt ergibt sich damit $|\lim_{n \rightarrow \infty} a_n| = \lim_{n \rightarrow \infty} |a_n|$, was nichts anderes als die Stetigkeit der Betragsfunktion bedeutet.

Wir können nun, analog zu Definition 4.3.1, den Begriff der reellen Cauchyfolge einführen. Bis auf weiteres meinen wir, sofern nichts anderes angegeben ist, mit dem Ausdruck $\varepsilon > C(0)$ eine beliebige *reelle* Zahl größer als das Nullelement.

Definition 4.5.4. Sei (a_n) eine Folge reeller Zahlen. Dann heißt (a_n) Cauchyfolge, wenn

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

Definition 4.5.5. Eine Folge reeller Intervalle $(I_n) = ([a_n, b_n])$ heißt (reelle) Intervallschachtelung, falls $I_{n+1} \subseteq I_n$ für alle $n \in \mathbb{N}$ gilt, und $\lim_{n \rightarrow \infty} (b_n - a_n) = C(0)$ ist.

Betrachten wir die Propositionen zu rationalen Cauchyfolgen und Intervallschachtelungen, so können wir feststellen, dass wir darin nur mit Eigenschaften gearbeitet haben, die wir auch für die reellen Zahlen beweisen können. Somit lassen sich Lemma 4.2.3,

Proposition 4.2.4, Proposition 4.2.5, Proposition 4.3.2 und Proposition 4.3.3 auf die reellen Zahlen übertragen, indem wir lediglich die Notation in den Beweisen anpassen (etwa $\varepsilon \cdot C(2)^{-1}$ anstelle von $\frac{\varepsilon}{2}$) [36, S. 25].

Sehr wichtig für unsere weiteren Überlegungen ist folgende Proposition über die Dichtigkeit der rationalen reellen Zahlen in den reellen Zahlen. Diese sagt aus, dass wir jede reelle Zahl, auch irrationale reelle Zahlen, beliebig genau durch rationale reelle Zahlen approximieren können.

Proposition 4.5.6. *Seien $C(a_n) < C(b_n)$ zwei verschiedene reelle Zahlen. Dann gibt es eine rationale reelle Zahl $C(c)$ mit $C(a_n) < C(c) < C(b_n)$.*

Beweis. Da es sich bei (a_n) und (b_n) um Cauchyfolgen handelt, gibt es nach Proposition 4.3.3 Intervallschachtelungen $(I_n) = ([p_n, q_n])$ und $(J_n) = ([r_n, s_n])$, sodass (a_n) ab einem Index in jedem I_N , und (b_n) ab einem Index in jedem J_N enthalten ist.

Da $C(q_n) = C(a_n) < C(b_n) = C(r_n)$ ist, gibt es nach Proposition 4.4.14 einen Index $N \in \mathbb{N}$, sodass $q_n < r_n$ für alle $n \geq N$ gilt. Insbesondere ist $q_N < r_N$, weswegen wir eine rationale Zahl c finden können, die $q_N < c < r_N$ erfüllt (etwa das arithmetische Mittel von q_N und r_N).

Es ist dann $C((c)_{n \in \mathbb{N}})$ eine rationale reelle Zahl. Zudem handelt es sich bei (q_n) nach Lemma 4.2.3 um eine monoton fallende Folge. Daher gilt

$$0 < c - q_N \leq c - q_{N+1} \leq c - q_{N+2} \leq \dots,$$

weswegen $C((c - q_n))$ eine positive reelle Zahl ist. Somit ist $C(c) = C(q_n) + C((c - q_n))$, also $C(q_n) < C(c)$. Ähnlich lässt sich zeigen, dass $C(c) < C(r_n)$ ist. $\square$

Um das Hauptresultat dieses Abschnitts beweisen zu können, benötigen wir folgendes Lemma. In der Formulierung und im Beweis ist die Unterscheidung zwischen konstanten und im Allgemeinen nichtkonstanten Folgen wichtig. Daher werden konstante Folgen durch die Verwendung eines unabhängigen Laufindex gekennzeichnet: $(a_n)_{k \in \mathbb{N}} = (a_n, a_n, a_n, \dots)$.

Lemma 4.5.7. *Sei (q_n) eine rationale Cauchyfolge. Dann ist $\lim_{n \rightarrow \infty} C((q_n)_{k \in \mathbb{N}}) = C(q_k)$.*

Beweis. Sei $\varepsilon > C(0)$. Wegen Proposition 4.5.6 gibt es ein $\eta \in \mathbb{Q}^+$ mit $C(0) < C(\eta) < \varepsilon$. Da es sich bei (q_n) um eine Cauchyfolge handelt, gibt es ein $N \in \mathbb{N}$, sodass $|q_m - q_n| < \eta$ für alle $n, m \geq N$ ist. Halten wir $n \geq N$ fest, so sind die Glieder von $(|q_k - q_n|)_{k \in \mathbb{N}}$ ab dem Index N kleiner als die Glieder der konstanten Folge $(\eta)_{k \in \mathbb{N}}$.

Aus Proposition 4.4.14 folgt damit, dass $C((|q_k - q_n|)_{k \in \mathbb{N}}) \leq C(\eta)$ ist. Mit Hilfe von Proposition 4.5.3 erhalten wir

$$\varepsilon > C(\eta) \geq C((|q_k - q_n|)_{k \in \mathbb{N}}) = |C((q_k - q_n)_{k \in \mathbb{N}})| = |C(q_k) - C((q_n)_{k \in \mathbb{N}})|.$$

Zusammengefasst: Für ein beliebiges $\varepsilon > 0$ gibt es einen Index $N \in \mathbb{N}$, sodass

$$|C(q_k) - C((q_n)_{k \in \mathbb{N}})| < \varepsilon$$

für $n \geq N$ gilt. Somit ist $\lim_{n \rightarrow \infty} C((q_n)_{k \in \mathbb{N}}) = C(q_k)$. □

Satz 4.5.8. *Die reellen Zahlen sind vollständig: Jede reelle Cauchyfolge ist konvergent.*

Beweis. Sei (b_n) eine reelle Cauchyfolge. Dann gibt es nach Proposition 4.3.3 (auf die reellen Zahlen übertragen) eine reelle Intervallschachtelung $(I_n) = ([a_n, c_n])$, sodass (b_n) ab einem Index in jedem Intervall enthalten ist.

Für $N \in \mathbb{N}$ finden wir wegen Proposition 4.5.6 in I_N sicher eine rationale reelle Zahl $C((q_N)_{k \in \mathbb{N}})$. Die Folge $(C((q_n)_{k \in \mathbb{N}}))_{n \in \mathbb{N}}$ jener rationalen reellen Zahlen ist eine reelle Cauchyfolge, da sie in einer Intervallschachtelung liegt. Für $\varepsilon \in \mathbb{Q}^+$ gibt es daher ein $N \in \mathbb{N}$, sodass

$$|C((q_m)_{k \in \mathbb{N}}) - C((q_n)_{k \in \mathbb{N}})| = |C((q_m - q_n)_{k \in \mathbb{N}})| < C(\varepsilon)$$

für alle $n, m \geq N$ erfüllt ist. Wegen Proposition 4.5.3 ist daher

$$C((|q_m - q_n|)_{k \in \mathbb{N}}) < C(\varepsilon),$$

und nach Proposition 4.4.14 bedeutet das, dass $|q_m - q_n| < \varepsilon$ ist. Es handelt sich somit bei $(q_n)_{n \in \mathbb{N}}$ um eine rationale Cauchyfolge.

Setzen wir $\beta := C((q_n)_{n \in \mathbb{N}})$. Wegen Lemma 4.5.7 ist $\lim_{n \rightarrow \infty} C((q_n)_{k \in \mathbb{N}}) = \beta$. Wir müssen nun zeigen, dass β in allen Intervallen enthalten ist.

Für $N \in \mathbb{N}$ ist $C((q_N)_{k \in \mathbb{N}}) \in I_N$, und da $I_n \subseteq I_N$ für alle $n \geq N$ gilt, ist $C((q_n)_{k \in \mathbb{N}}) \in I_N$ für alle $n \geq N$. Wegen der Eigenschaften von Grenzwerten muss daher $\beta \in I_N$ sein. Die reelle Zahl β liegt somit in allen Intervallen, weswegen $\lim_{n \rightarrow \infty} b_n = \beta$ ist. $\square$

Als erste Anwendung der Vollständigkeit beweisen wir, dass Gleichungen ähnlich zu der aus Abschnitt 4.2 in den reellen Zahlen lösbar sind. Zur besseren Lesbarkeit verwenden wir die Notation $\frac{a}{b} = a \cdot b^{-1}$, und schreiben $a = C(a)$ für rationale reelle Zahlen; diese Schreibweise wird im nächsten Abschnitt formalisiert.

Proposition 4.5.9. *Sei $a \in \mathbb{R}$ und $k \in \mathbb{N}$ mit $k \geq 1$. Dann ist die Gleichung $x^k = a$ in den reellen Zahlen stets lösbar, wenn k ungerade ist. Falls k gerade ist, so ist die Gleichung lösbar, wenn a nichtnegativ ist.*

Beweis. Für $a = 0$ ist die Lösung durch $x = 0$ gegeben. Andernfalls gehen wir ähnlich vor wie in Abschnitt 4.2. Zunächst definieren wir das Intervall $I_0 = [a_0, b_0]$ mit $a_0 := \min\{-1, -|a|\}$ und $b_0 := \max\{1, |a|\}$. Falls $|a| \geq 1$ ist, so muss wegen $|a|^k \geq |a|$ das $x \in \mathbb{R}$, welches diese Gleichung erfüllt in diesem Intervall liegen – vorausgesetzt es existiert. Für $|a| < 1$ ist $|a|^k \leq |a|$, und damit ebenfalls $x \in I_0$.

Die nachfolgenden Intervalle definieren wir wie folgt:

$$I_{n+1} = [a_{n+1}, b_{n+1}] := \begin{cases} \left[a_n, \frac{a_n + b_n}{2} \right] & \text{wenn } \left(\frac{a_n + b_n}{2} \right)^k \geq a, \\ \left[\frac{a_n + b_n}{2}, b_n \right] & \text{wenn } \left(\frac{a_n + b_n}{2} \right)^k < a. \end{cases}$$

Per Konstruktion ist $a_n^k \leq a$ und $b_n^k \geq a$ für alle $n \in \mathbb{N}$. Wegen $b_n - a_n = \left(\frac{1}{2}\right)^n \cdot 2 \cdot b_0$ und $\lim_{n \rightarrow \infty} \left(\frac{1}{2}\right)^n = 0$ handelt es sich bei (I_n) um eine reelle Intervallschachtelung. Daher sind (a_n) und (b_n) Cauchyfolgen, die nach Satz 4.5.8 gegen eine Zahl $\alpha \in \mathbb{R}$ konvergieren. Aufgrund der Grenzwerteigenschaften ist $\alpha^k \leq a$ und $\alpha^k \geq a$, woraus $\alpha^k = a$ folgt. Somit ist α eine Lösung der Gleichung $x^k = a$.

Ist $a \in \mathbb{R}^+$, so ist $\left(\frac{-a+a}{2}\right)^k = 0 < a$, woraus $a_1 = 0$ folgt. Insbesondere ist damit $\alpha \geq a_0 = 0$ nichtnegativ. $\square$

Das obige Resultat können wir dazu verwenden, die Wurzelfunktion auf den nichtnegativen reellen Zahlen zu definieren.

Definition 4.5.10. Sei $a \in \mathbb{R} \setminus \mathbb{R}^-$ und $n \in \mathbb{N}$ mit $n \geq 1$. Dann definieren wir $\sqrt[n]{a}$ als die nichtnegative Lösung der Gleichung $x^n = a$.

Für das Rechnen mit Wurzelausdrücken gelten die bekannten Rechenregeln.

4.6 Einbettung von $\mathbb{Q}$ in $\mathbb{R}$

Da wir die reellen Zahlen über Cauchyfolgen definiert haben, mussten wir bisher mit einer etwas unintuitiven Notation arbeiten – sind wir es doch eher gewohnt, dass sich reelle Zahlen als Ergebnisse wie $\sqrt{2}$ oder als Dezimalzahlen wie 1,414213562 ... anschreiben lassen. In diesem Abschnitt geht es daher darum, die rationalen Zahlen in die reellen Zahlen einzubetten, und die bekannte Notation „wiederherzustellen“.

Eine Einbettung im Sinne von Definition 2.4.1 zu finden ist nicht sonderlich schwer: Wir ordnen jeder rationalen Zahl $a \in \mathbb{Q}$ die Äquivalenzklasse der konstanten Folge $(a)_{n \in \mathbb{N}}$ zu. Diese ist offensichtlich eine Cauchyfolge.

Proposition 4.6.1. Sei $\varphi : \mathbb{Q} \rightarrow \mathbb{R}$ mit $\varphi(a) := C((a))$. Dann handelt es sich bei φ um eine Einbettung.

Beweis. Seien $a, b \in \mathbb{Q}$. Dann ist

$$\varphi(a + b) = C((a + b)) = C((a) + (b)) = C((a)) + C((b)) = \varphi(a) + \varphi(b).$$

Zudem ist

$$\varphi(a \cdot b) = C((a \cdot b)) = C((a) \cdot (b)) = C((a)) \cdot C((b)) = \varphi(a) \cdot \varphi(b).$$

Für den Beweis der Injektivität nehmen wir an, dass $\varphi(a) = C((a)) = C((b)) = \varphi(b)$ ist. Dann ist nach Definition 4.3.6 $\lim_{n \rightarrow \infty} (a - b) = 0$. Durch Anwenden der Rechenregeln für Grenzwerte und Umformen des Ausdrucks erhalten wir daraus direkt $a = b$. $\square$

Auch die Ordnung der rationalen Zahlen ist mit der Ordnung der reellen Zahlen verträglich.

Proposition 4.6.2. *Seien $a, b \in \mathbb{Q}$. Dann ist $a \leq b$ genau dann, wenn $\varphi(a) \leq \varphi(b)$ ist.*

Beweis. Angenommen, es wäre $a \leq b$. Dann gibt es ein $c \in \mathbb{Q} \setminus \mathbb{Q}^-$ mit $b = a + c$. Es ist $\varphi(b) = \varphi(a + c) = \varphi(a) + \varphi(c)$. Da c eine nichtnegative rationale Zahl ist, handelt es sich bei $C((c))$ um eine nichtnegative reelle Zahl. Damit ist $\varphi(a) \leq \varphi(b)$.

Um die umgekehrte Richtung zu beweisen, treffen wir eine Fallunterscheidung. Ist $\varphi(a) = \varphi(b)$, so gilt $\lim_{n \rightarrow \infty} (b - a) = 0$, woraus direkt $a = b$ folgt. Ist hingegen, ohne Beschränkung der Allgemeinheit, $\varphi(a) < \varphi(b)$, so ist $C((b - a))$ eine positive reelle Zahl. Nach Definition 4.4.6 gibt es daher ein $c \in \mathbb{Q}^+$ mit $b - a \geq c$. Damit ist $b \geq a + c > a$, also $a < b$. $\square$

Wir haben somit gezeigt, dass die rationalen Zahlen mit den reellen Zahlen verträglich sind, und schreiben $\mathbb{Q} \subseteq \mathbb{R}$ um auszudrücken, dass die reellen Zahlen eine Erweiterung der rationalen Zahlen sind.

Für die Notation einigen wir uns zunächst darauf, jede (rationale) reelle Zahl $C((a))$ mit $a \in \mathbb{Q}$ künftig als a anzuschreiben. Außerdem benutzen wir die bekannte Bruchschreibweise $\frac{\alpha}{\beta}$ mit $\beta \neq 0$ anstelle von $\alpha \cdot \beta^{-1}$ für $\alpha, \beta \in \mathbb{R}$, wobei $\beta \neq 0$.

Von Interesse sind nun die irrationalen Zahlen, sprich die Äquivalenzklassen jener rationalen Cauchyfolgen, die in $\mathbb{Q}$ keinen Grenzwert besitzen. Wir zeigen, dass wir jede reelle Zahl beliebig genau durch eine rationale Dezimalzahl approximieren können. Zu jeder reellen Zahl lassen sich damit beliebig viele Nachkommastellen ausrechnen.

Proposition 4.6.3. *Sei $\alpha \in \mathbb{R}$ und $\varepsilon > 0$. Dann gibt es ein $a \in \mathbb{Q}$, sodass $|\alpha - a| < \varepsilon$ ist.*

Beweis. Betrachten wir die reellen Zahlen α und $\alpha + \varepsilon$. Nach Proposition 4.5.6 gibt es eine rationale Zahl $a \in \mathbb{Q}$ mit $\alpha < a < \alpha + \varepsilon$. Es ist daher $|\alpha - a| < \varepsilon$. $\square$

Zudem können wir Proposition 3.3.4 ergänzen: Ist $a_0, a_1 a_2 a_3 \dots$ eine Dezimaldarstellung, so entspricht dieser sicher eine reelle Zahl.

Proposition 4.6.4. Sei $a_0, a_1 a_2 a_3 \dots$ eine Dezimaldarstellung. Dann ist

$$\alpha = a_0 + \sum_{n=1}^{\infty} a_n \cdot \left(\frac{1}{10}\right)^n$$

eine reelle Zahl.

Beweis. Es ist $0 \leq a_n \leq 9$ für alle a_n mit $n \geq 1$. Zudem konvergiert $\sum_{n=1}^{\infty} 9 \cdot \left(\frac{1}{10}\right)^n$, da es sich um eine geometrische Reihe handelt, und $\frac{1}{10} < 1$ ist. Wegen $a_n \cdot \left(\frac{1}{10}\right)^n \leq 9 \cdot \left(\frac{1}{10}\right)^n$ konvergiert die ursprüngliche Reihe nach dem Majorantenkriterium. $\square$

Da irrationale Zahlen nach Proposition 3.3.4 eine nicht abbrechende, nicht periodische Dezimaldarstellung besitzen, sind wir bei der Darstellung irrationaler Zahlen in der Praxis an endliche Genauigkeit gebunden. Sie ist daher vor allem für numerische Anwendungen interessant. In der Regel ist es sinnvoller, einer reellen Zahl eine eindeutige Eigenschaft zuzuweisen, und die Zahl anschließend konkret zu benennen (etwa π oder e), oder sie als Bild einer Funktion darzustellen, wie $\sqrt{2}$.

Durch die Erweiterung auf die reellen Zahlen haben wir gewonnen, dass der Zahlenstrahl keine Lücken enthält: Jede Cauchyfolge ist konvergent. Dennoch sind wir nach wie vor nicht dazu in der Lage, die Nullstellen beliebiger Polynome zu finden. So gibt es etwa keine reelle Zahl mit $x^2 = -1$, denn sowohl das Produkt zweier positiver, wie auch zweier negativer Zahlen ist stets positiv. Die Suche nach Nullstellen von Polynomen wird uns im nächsten Kapitel zu den komplexen Zahlen beschäftigen.

4.7 Unendlichkeit der Zahlenmengen

Das Kapitel zu den reellen Zahlen bietet sich gut dafür an, die Größe der bisher eingeführten Zahlenmengen zu behandeln. Intuitiv fällt es uns nicht schwer zu sagen, dass alle bisher bekannten Zahlenmengen „unendlich groß“ sind. Wie dieser Abschnitt zeigen wird, versagt unsere Intuition jedoch gerade bei unendlich großen Zahlenmengen – es ist daher notwendig, dass wir uns diesem Begriff behutsam nähern, und sorgfältig damit arbeiten.

Um die Größe von Mengen untersuchen zu können, benötigen wir zunächst die Möglichkeit, die Größen in Relation zueinander zu setzen. Dazu verwenden wir die Idee der paarweisen Zuordnung, die bereits viele tausend Jahre alt ist, wie wir in Abschnitt 1.1 gezeigt haben.

Definition 4.7.1 ([19, S. 65]). Seien A und B Mengen. Dann definieren wir:

- (1) $|A| \leq |B|$, wenn es eine injektive Funktion $f : A \rightarrow B$ gibt,
- (2) $|A| = |B|$, wenn es eine bijektive Funktion $f : A \rightarrow B$ gibt.

Im ersten Fall sagen wir, dass die Mächtigkeit von A kleiner gleich der Mächtigkeit von B ist. Ist der zweite Punkt erfüllt, so nennen wir A und B gleich mächtig: A und B haben gleich viele Elemente.

An dieser Stelle können wir nicht formal beweisen, dass es sich bei $\leq$ um eine Ordnungsrelation, und bei $=$ um eine Äquivalenzrelation handelt, denn wir haben nicht definiert, um was für ein Objekt es sich bei $|A|$ überhaupt handelt. Die Antwort auf diese Frage führt auf das Gebiet der Kardinalzahlen, und würde leider den Umfang dieser Diplomarbeit sprengen [vgl. 19, S. 162–182]. Für unsere Zwecke ist es ausreichend, wenn wir weniger formal vorgehen, und für eine endliche Menge mit $|A|$ die Anzahl der Elemente von A meinen, während wir für einige ausgewählte unendliche Mengen spezielle Bezeichnungen der Kardinalitäten einführen.

Zunächst benötigen wir jedoch eine Definition von „unendlich“. Als Ausgangspunkt benutzen wir die Menge der natürlichen Zahlen, die uns intuitiv als unendliche Menge erscheint. Können wir eine surjektive Zuordnung einer Menge A zur Menge der natürlichen Zahlen finden, so enthält A sicherlich mindestens so viele Elemente wie $\mathbb{N}$, und wäre unserer Intuition nach ebenfalls unendlich.

Definition 4.7.2 ([19, S. 104]). Wir nennen eine Menge A unendlich, falls es eine surjektive Funktion $f : A \rightarrow \mathbb{N}$ gibt. Ist eine Menge nicht unendlich, so nennen wir sie endlich.

Aus dieser Definition folgt, dass die Menge der natürlichen Zahlen die kleinste unendliche Menge ist.

Proposition 4.7.3. *Sei A eine unendliche Menge. Dann ist $|\mathbb{N}| \leq |A|$.*

Beweis. Da A unendlich ist, gibt es eine surjektive Funktion $f : A \rightarrow \mathbb{N}$. Für alle $n \in \mathbb{N}$ ist daher das Urbild $f^{-1}(n) = \{a \in A \mid f(a) = n\}$ nichtleer, weswegen wir nach dem Auswahlaxiom genau ein Element a_n aus dieser Menge auswählen können. Definieren wir nun $g : \mathbb{N} \rightarrow A$ durch $g(n) := a_n$, so folgt aus $g(n) = g(m)$, dass a_n und a_m im Urbild $f^{-1}(n)$ liegen, womit $n = m$ ist [33]. Somit ist g injektiv, und $|\mathbb{N}| \leq |A|$ nach Definition 4.7.1. $\square$

Die uns bekannten Zahlenmengen sind, wie erwartet, unendlich. Zum Beweis betrachten wir einfach eine Funktion, die alle natürlichen Zahlen auf sich selbst abbildet, und alle übrigen (ganzen/rationalen/reellen) Zahlen auf das Nullelement.

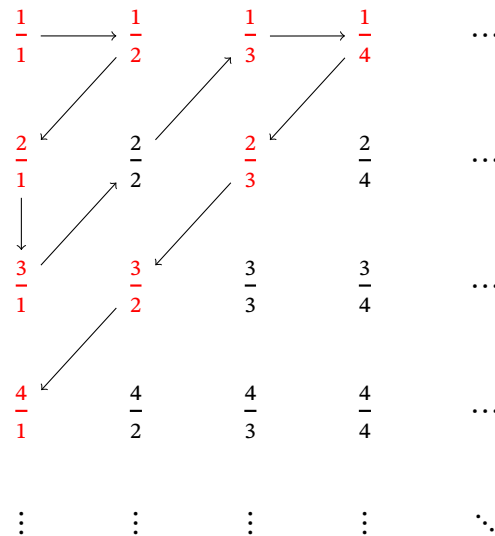
Eine interessante Eigenschaft unendlicher Mengen lässt sich beobachten, wenn wir etwa die Menge $N = \{n \in \mathbb{N} \mid 2 \mid n\}$ und die Funktion $f : \mathbb{N} \rightarrow N$ mit $f(n) = 2 \cdot n$ untersuchen. Wegen der Definition der Menge handelt es sich bei f sicher um eine Surjektion, und aus $2 \cdot n = f(n) = f(m) = 2 \cdot m$ folgt mit der Kürzungsregel, dass $n = m$ ist. Die Funktion f ist daher auch injektiv, und damit insgesamt bijektiv. Nach Definition 4.7.1 ist daher $|N| = |\mathbb{N}|$. Die Mengen haben also gleich viele Elemente, obwohl N eine echte Teilmenge der natürlichen Zahlen ist!

Diese Besonderheit widerspricht unserer Intuition, dass ein Teil, dem etwas auf das Ganze fehlt, stets kleiner als jenes Ganze ist. Es lässt sich zeigen, dass diese Eigenschaft charakteristisch für unendliche Mengen ist, und sogar eine äquivalente Definition des Unendlichkeitsbegriffs ermöglicht [19, S. 93, 104].

Auch die uns bekannten Zahlenmengen bleiben von solchen Besonderheiten nicht verschont, wie folgende Proposition zeigt.

Proposition 4.7.4. *Es ist $|\mathbb{N}| = |\mathbb{Z}| = |\mathbb{Q}|$.*

Beweis. Betrachten wir die Auflistung $0, 1, -1, 2, -2, 3, -3, \dots$. In dieser Liste kommen offenbar alle ganzen Zahlen vor. Gehen wir sie der Reihe nach durch (nulltes Element, erstes Element, ...), so erhalten wir dadurch eine bijektive Abbildung von den natürlichen Zahlen auf die ganzen Zahlen. Es ist somit $|\mathbb{N}| = |\mathbb{Z}|$.



Wir erhalten somit die Aufzählung $1, \frac{1}{2}, 2, 3, \frac{1}{3}, \frac{1}{4}, \dots$. Zur Erweiterung auf die nichtpositiven rationalen Zahlen fügen wir an den Beginn der Liste die Null an, und nach jedem Element dessen additives Inverses: $0, 1, -1, \frac{1}{2}, -\frac{1}{2}, 2, -2, 3, -3, \frac{1}{3}, -\frac{1}{3}, \frac{1}{4}, -\frac{1}{4}, \dots$. Gehen wir diese Auflistung nun der Reihe nach durch, erhalten wir eine Bijektion von $\mathbb{N}$ nach $\mathbb{Q}$. Es ist somit $|\mathbb{N}| = |\mathbb{Q}|$. $\square$

Definition 4.7.5 ([19, S. 109]). Eine Menge A heißt abzählbar unendlich, falls es eine bijektive Funktion $f : \mathbb{N} \rightarrow A$ gibt. Für die Kardinalität der natürlichen Zahlen schreiben wir $|\mathbb{N}| = \aleph_0$.

Die ganzen und rationalen Zahlen sind somit abzählbar unendlich, und es ist $|\mathbb{Z}| = |\mathbb{Q}| = \aleph_0$.

Es bleibt nun die Frage, ob auch die reellen Zahlen abzählbar unendlich sind. Die Antwort auf diese Frage hat Georg Cantor im Jahr 1873 gefunden, wobei vor allem sein späterer Beweis aus dem Jahr 1891 heute weit bekannt ist [19, S. 121].

Proposition 4.7.6 ([19, S. 121]). *Die reellen Zahlen sind überabzählbar, es gibt keine surjektive Funktion $f : \mathbb{N} \rightarrow \mathbb{R}$.*

Beweis. Wir zeigen, dass es bereits keine Abzählung des Intervalls $[0,1] \subseteq \mathbb{R}$ geben kann. Nehmen wir dazu an, es gäbe eine solche, also eine surjektive Funktion $f : \mathbb{N} \rightarrow [0,1]$:

$$\begin{aligned} f(0) &= a_{0,0}, a_{0,1}a_{0,2}a_{0,3}a_{0,4}a_{0,5} \dots \\ f(1) &= a_{1,0}, a_{1,1}a_{1,2}a_{1,3}a_{1,4}a_{1,5} \dots \\ f(2) &= a_{2,0}, a_{2,1}a_{2,2}a_{2,3}a_{2,4}a_{2,5} \dots \\ f(3) &= a_{3,0}, a_{3,1}a_{3,2}a_{3,3}a_{3,4}a_{3,5} \dots \\ f(4) &= a_{4,0}, a_{4,1}a_{4,2}a_{4,3}a_{4,4}a_{4,5} \dots \\ &\vdots \end{aligned}$$

Wir betrachten nun die Zahl $\alpha := 0, b_1b_2b_3b_4b_5 \dots$, wobei b_n wie folgt definiert ist.

$$b_n := \begin{cases} 1 & \text{wenn } a_{n,n} = 2, \\ 2 & \text{wenn } a_{n,n} \neq 2. \end{cases}$$

Dann ist jedenfalls $\alpha \neq f(0)$, da die erste Nachkommastelle verschieden ist. Es ist auch $\alpha \neq f(1)$, da die zweite Nachkommastelle verschieden ist. Allgemein ist $\alpha \neq f(n)$ für alle $n \in \mathbb{N}$, da sich die n -te Nachkommastelle unterscheidet. Wir haben somit eine Zahl $\alpha \in [0,1]$ gefunden, die nicht im Bild der Surjektion liegt – ein Widerspruch.

Daher kann es keine Abzählung des Intervalls $[0,1]$, und somit auch der reellen Zahlen geben. $\square$

Die Mächtigkeit der reellen Zahlen wird in der Regel mit $|\mathbb{R}| = \mathfrak{c}$ bezeichnet, wobei $\mathfrak{c}$ an „Kontinuum“ erinnern soll. Da sich zeigen lässt, dass $|\mathcal{P}(\mathbb{N})| = \mathfrak{c}$ ist, und für die

Größe der Potenzmenge einer endlichen Menge A stets $|\mathcal{P}(A)| = 2^{|A|}$ gilt [19, S. 134–135], schreibt man manchmal auch $\mathfrak{c} = 2^{\aleph_0}$.

Für die zukünftigen Erweiterungen auf die komplexen Zahlen, die Quaternionen und Oktonionen kann man beweisen, dass diese gleich mächtig zu den reellen Zahlen sind.

Proposition 4.7.7 ([19, S. 130–131]). *Für $n \in \mathbb{N}$ mit $n \geq 1$ ist $|\mathbb{R}^n| = \mathfrak{c}$.*

Wir haben bewiesen, dass die reellen Zahlen mächtiger als die Menge der natürlichen Zahlen sind. Die Frage ist nun, ob es Mengen gibt, deren Mächtigkeit zwischen $\aleph_0$ und $\mathfrak{c}$ liegt? Georg Cantor hat bereits im Jahr 1878 vermutet, dass dies nicht der Fall ist, konnte seine Vermutung allerdings nicht beweisen, und hat „die genaue Untersuchung dieser Frage [...] auf eine spätere Gelegenheit“ verschoben [9, S. 258]. David Hilbert setzte die Lösung dieser mittlerweile als *Kontinuumshypothese* bekannten Behauptung anlässlich seines Vortrags auf dem internationalen Mathematiker-Kongress in Paris im Jahr 1900 auf seine Liste der 23 Probleme, die für das 20. Jahrhundert von Bedeutung wären [28, S. 263–264].

Im Jahr 1938 konnte Kurt Gödel zeigen, dass die Kontinuumshypothese mit ZFC nicht widerlegbar ist. Es dauerte nochmals knappe dreißig Jahre, bis Paul Cohen 1963 zeigen konnte, dass sie sich mit ZFC auch nicht beweisen lässt [19, S. 152]. Die Kontinuumshypothese ist somit eine jener unentscheidbaren Aussagen, deren Existenz aus dem ersten Gödelschen Unvollständigkeitssatz folgt [24, S. 193]. Wollen wir daher die Frage um die Gültigkeit der Kontinuumshypothese beantworten, müssen wir ZFC um ein Axiom erweitern, das sie bestätigt oder verneint [19, S. 151].

5 Die komplexen Zahlen

Die Erweiterung auf die reellen Zahlen war dadurch motiviert, die Lücken im Zahlenstrahl zu schließen. Wir haben bereits festgestellt, dass wir nun zwar die Gleichung $x^2 = 2$ lösen können, es allerdings für die Gleichung $x^2 = -1$ keine Lösung in den reellen Zahlen gibt, da Quadratzahlen stets positiv sind.

Durch die Erweiterung auf die komplexen Zahlen wollen wir diese algebraische Lücke schließen, was uns auch gelingen wird: Der Fundamentalsatz der Algebra besagt, dass jedes nichtkonstante Polynom über den komplexen Zahlen mindestens eine Nullstelle besitzt.

Obwohl der Name „komplexe“ Zahlen vermuten lässt, dass es sich um einen besonders schwierigen Zahlenbereich handelt, werden wir feststellen, dass sich die komplexen Zahlen wesentlich einfacher als alle bisherigen Zahlenbereiche definieren lassen (tatsächlich ist „komplex“ im Sinne von „zusammengesetzt“ zu verstehen). Für die Lösbarkeit von Gleichungen wie $x^2 = -1$ bezahlen wir aber einen Preis. Wie wir sehen werden, lässt sich die Existenz einer negativen Quadratzahl nicht in Einklang mit den Ordnungsaxiomen bringen. Wir gewinnen zwar die algebraische Abgeschlossenheit, verlieren dafür aber die uns gewohnte Anordnung der Zahlen.

Den komplexen Zahlen haftete lange Zeit ein Hauch von Mystik an; sie wurden von vielen nicht als „echte Zahlen“ anerkannt, was sich unter anderem in der Bezeichnung „*imaginäre* Einheit“ widerspiegelt. Wenn wir im Titel dieser Arbeit von „eingebildeten Zahlen“ sprechen, meinen wir daher die komplexen Zahlen.

5.1 Lösungen von Gleichungen höheren Grades

Wie die Einleitung dieses Kapitels andeutet, traten komplexen Zahlen historisch gesehen erstmals auf der Suche nach Lösungen für Gleichungen höheren Grades auf. Bereits seit der Antike ist bekannt, wie sich quadratische Gleichungen der Form $x^2 + p \cdot x + q = 0$ lösen lassen: Wir können diese Gleichung durch Addieren von $\frac{p^2}{4} - q$ auf ein vollständiges Quadrat ergänzen, und erhalten dadurch

$$\left(x + \frac{p}{2}\right)^2 = \frac{p^2}{4} - q. \quad (5.1)$$

Ist $\frac{p^2}{4} - q \geq 0$, so können wir die Wurzel ziehen, und erhalten die Lösungen

$$x_{1,2} = -\frac{p}{2} \pm \sqrt{\frac{p^2}{4} - q}.$$

Für dieses Kapitel ist vor allem der Fall $\frac{p^2}{4} - q < 0$ interessant. Da Quadrate reeller Zahlen stets positiv sind, gibt es in dem Fall keine reelle Zahl, welche Gleichung 5.1 erfüllt.

Die Suche nach einer Lösungsformel für Gleichungen dritten Grades hat vor allem die Mathematiker des 15. und 16. Jahrhunderts beschäftigt, allen voran die Italiener Scipione del Ferro, Niccolò Tartaglia und Gerolamo Cardano [5, S. 317–319]. Auf Cardano geht ein Lösungsweg für die allgemeine kubische Gleichung zurück, welcher auf der Arbeit Tartaglias und del Ferros basiert [5, S. 320; 39, S. 16]. Er gab auch eine Methode an, um Gleichungen vierter Ordnung auf Gleichungen dritter Ordnung zurückzuführen [56, S. 392]. Eine Lösungsformel für Gleichungen fünfter Ordnung und höher konnten er und nachfolgende Generationen von Mathematikern nicht finden – was daran liegt, dass es diese nicht gibt, wie Niels Henrik Abel beweisen konnte [12, S. 287].

Cardano war einer der ersten, die komplexe Zahlen als Wurzeln negativer reeller Zahlen verwendeten, wenn auch nur „mit Kopfschmerzen“ [39, S. 17]. In seinem Werk, der *Ars Magna*, stellte er die Aufgabe, die Zahl Zehn derart in zwei Teile zu teilen, dass das

Produkt der Teile gleich Vierzig ist. Algebraisch formuliert: Finde $x, y \in \mathbb{R}$ mit

$$x + y = 10,$$

$$x \cdot y = 40.$$

Einsetzen der ersten Gleichung in die zweite liefert $x \cdot (10 - x) = 40$, beziehungsweise $x^2 - 10 \cdot x + 40 = 0$. Nach der obigen Lösungsformel würde die Gleichung die beiden Lösungen $x_{1,2} = 5 \pm \sqrt{-15}$ besitzen, allerdings ist -15 eine negative Zahl, und besitzt demnach keine reelle Wurzel. Interessant ist nun, dass Cardano vorübergehend $\sqrt{-15}$ als Zahl akzeptierte, mit der man wie gewohnt rechnen kann. Wir erhalten dann

$$x_1 + x_2 = (5 + \sqrt{-15}) + (5 - \sqrt{-15}) = 10,$$

$$x_1 \cdot x_2 = (5 + \sqrt{-15}) \cdot (5 - \sqrt{-15}) = 25 - (\sqrt{-15})^2 = 40.$$

Die „Zahlen“ $x_{1,2}$ sind daher eine Lösung der Aufgabe, obwohl sie sich in keinen bisher bekannten Zahlenbereich einordnen lassen [39, S. 17].

Eine weitere bizarre Situation erhalten wir, wenn wir Cardanos Lösungsweg ([vgl. 12, S. 29–31]) auf die Gleichung

$$x^3 - 2 \cdot x^2 - 5 \cdot x + 6 = 0 \tag{5.2}$$

anwenden. Durch die Substitution $x = z + \frac{2}{3}$ erhalten wir

$$z^3 - \frac{19}{3} \cdot z + \frac{56}{27} = 0,$$

also eine kubische Gleichung der Form $x^3 + p \cdot x + q = 0$, die Cardano dank der Arbeit Tartaglias lösen konnte. Wenden wir die Cardanischen Formeln [12, S. 30–31] darauf an, so erhalten wir eine Lösung a als

$$a = \sqrt[3]{-\frac{28}{27} + \sqrt{\frac{784}{729} - \frac{6859}{729}}} + \sqrt[3]{-\frac{28}{27} - \sqrt{\frac{784}{729} - \frac{6859}{729}}}. \tag{5.3}$$

Da Gleichung 5.2 die Nullstellen $-2, 1$ und 3 besitzt, und es keine weiteren Lösungen der Gleichung geben kann (siehe Proposition 5.6.7), muss a zurücksubstituiert einer dieser

drei Zahlen entsprechen – obwohl sie eine Darstellung besitzt, in der Quadratwurzeln negativer Zahlen vorkommen, und anschließend die Kubikwurzel dieses Ergebnisses verwendet wird. Cardano wusste nicht, wie man in diesem Fall vorgehen soll, und nannte dies den *Casus irreducibilis*, der immer im Fall dreier reeller Lösungen auftritt [12, S. 32].

Der Mathematiker Rafael Bombelli führte Cardanos Überlegungen weiter, und war einer der ersten, welche die Existenz komplexer Zahlen akzeptierten. Bombelli hatte die Idee, die Summanden in der Cardanischen Formel (Gleichung 5.3) als konjugiert komplexe Zahlen anzusehen, die sich nur um ein Vorzeichen unterscheiden.

$$\begin{aligned}\sqrt[3]{-\frac{28}{27} + \sqrt{\frac{784}{729} - \frac{6859}{729}}} &= a + b \cdot \sqrt{-1} \\ \sqrt[3]{-\frac{28}{27} - \sqrt{\frac{784}{729} - \frac{6859}{729}}} &= a - b \cdot \sqrt{-1}\end{aligned}$$

Durch Ausrechnen der Werte für a und b konnte er zeigen, dass die Summe dieser komplexen Zahlen tatsächlich eine reelle Zahl ist [39, S. 19–22; 5, S. 324–326].

Bis zur generellen Akzeptanz komplexer Zahlen dauerte es allerdings bis ins 19. Jahrhundert. Vor allem die Bemühungen von Carl Friedrich Gauß sind hervorzuheben, der zu dem Thema wesentliche Beiträge geleistet hat [39, S. 82].

5.2 Konstruktion der komplexen Zahlen

Gauß hatte bereits im Jahr 1796 die Idee, komplexe Zahlen als Paare von reellen Zahlen anzusehen, die man in einem kartesischen Koordinatensystem darstellen kann [39, S. 82]. Dies wurde 1835 von William Hamilton, dem Entdecker der Quaternionen, formalisiert, und stellt die Basis für unsere Konstruktion der komplexen Zahlen dar [20, S. 52].

Definition 5.2.1. Die Menge der komplexen Zahlen $\mathbb{C}$ ist definiert als

$$\mathbb{C} := \mathbb{R} \times \mathbb{R} = \{(a, b) \mid a, b \in \mathbb{R}\}.$$

Die komplexen Zahlen können wir nach dieser Definition auch als Vektorraum über den reellen Zahlen verstehen, welcher die Standardbasis $\{(1,0), (0,1)\}$ besitzt. Setzen wir nun $e := (1,0)$ und $i := (0,1)$, so können wir jede komplexe Zahl (λ_1, λ_2) bezüglich dieser Basis als

$$(\lambda_1, \lambda_2) = \lambda_1 \cdot e + \lambda_2 \cdot i$$

ausdrücken, wobei $\lambda_1, \lambda_2 \in \mathbb{R}$ sind. Wir nennen i auch die *imaginäre Einheit*.

Definition 5.2.2. Ist $(a, b) \in \mathbb{C}$, so definieren wir die konjugiert komplexe Zahl $\overline{(a, b)}$ als

$$\overline{(a, b)} := (a, -b).$$

Definition 5.2.3. Für $(a, b) \in \mathbb{C}$ definieren wir den Betrag dieser Zahl als

$$|(a, b)| := \sqrt{a^2 + b^2}.$$

Es lässt sich zeigen, dass die Betragsfunktion auf den komplexen Zahlen alle gewohnten Eigenschaften, unter anderem die Produktregel, aufweist; insbesondere erfüllt sie die Dreiecksungleichung.

5.3 Rechenoperationen auf $\mathbb{C}$

Bei den komplexen Zahlen handelt es sich, wie bei den reellen Zahlen, um einen vollständigen Körper. Anders als in den bisherigen Zahlenbereichen werden wir allerdings keinen Erfolg damit haben, eine mit den in diesem Abschnitt definierten Rechenoperationen verträgliche Ordnung zu finden.

Um die folgenden Definitionen zu motivieren, nehmen wir zunächst einige uns bekannte Tatsachen komplexer Zahlen, wie Rechengesetze und Schreibweisen, als gegeben an. Dann ist $(a + b \cdot i) + (c + d \cdot i) = (a + c) + (b + d) \cdot i$. Bei der Addition werden somit Real- und Imaginärteil komponentenweise addiert.

Definition 5.3.1. Die Addition $+$: $\mathbb{C} \times \mathbb{C} \rightarrow \mathbb{C}$ ist definiert durch

$$(a, b) + (c, d) := (a + c, b + d).$$

Für die Definition der Multiplikation betrachten wir

$$\begin{aligned}
 (a + b \cdot i) \cdot (c + d \cdot i) &= a \cdot c + a \cdot d \cdot i + b \cdot i \cdot c + b \cdot i \cdot d \cdot i = \\
 &= a \cdot c + b \cdot d \cdot \underbrace{i^2}_{=-1} + (a \cdot d + b \cdot c) \cdot i = \\
 &= a \cdot c - b \cdot d + (a \cdot d + b \cdot c) \cdot i.
 \end{aligned}$$

Wir definieren die Multiplikation daher wie folgt.

Definition 5.3.2. Die Multiplikation $\cdot : \mathbb{C} \times \mathbb{C} \rightarrow \mathbb{C}$ ist definiert durch

$$(a, b) \cdot (c, d) := (a \cdot c - b \cdot d, a \cdot d + b \cdot c).$$

Wegen $i = (0,1)$ ist $i^2 = (0,1) \cdot (0,1) = (-1,0)$, wobei wir sehen werden, dass dies der reellen Zahl -1 entspricht.

Proposition 5.3.3. $(\mathbb{C}, +, \cdot, (0,0), (1,0))$ ist ein Körper. Es sind daher die Körperaxiome erfüllt (vgl. [50, S. 264]):

- (1) $\forall a, b, c \in \mathbb{C} : (a + b) + c = a + (b + c)$ (Assoziativität von +).
- (2) $\forall a, b \in \mathbb{C} : a + b = b + a$ (Kommutativität von +).
- (3) $\forall a \in \mathbb{C} : a + (0,0) = a$ (Nullelement).
- (4) $\forall a \in \mathbb{C} : \exists -a \in \mathbb{C} : a + (-a) = (0,0)$ (Invertierbarkeit von +).
- (5) $\forall a, b, c \in \mathbb{C} : (a \cdot b) \cdot c = a \cdot (b \cdot c)$ (Assoziativität von $\cdot$).
- (6) $\forall a, b \in \mathbb{C} : a \cdot b = b \cdot a$ (Kommutativität von $\cdot$).
- (7) $\forall a \in \mathbb{C} : a \cdot (1,0) = a$ (Einselement).
- (8) $\forall a \in \mathbb{C} \setminus \{(0,0)\} : \exists a^{-1} \in \mathbb{C} : a \cdot a^{-1} = (1,0)$ (Invertierbarkeit von $\cdot$).
- (9) $\forall a, b, c \in \mathbb{C} : a \cdot (b + c) = a \cdot b + a \cdot c$ (Distributivität).

Beweis. Seien $a = (a_1, a_2)$, $b = (b_1, b_2)$ und $c = (c_1, c_2)$ komplexe Zahlen. Die Eigenschaften (1) bis (4) folgen direkt aus der Tatsache, dass die reellen Zahlen einen Körper bilden, und die Addition zweier komplexer Zahlen komponentenweise durchgeführt wird.

Zur Assoziativität der Multiplikation: Zwecks besserer Lesbarkeit schreiben wir komplexe Zahlen als Zeilenvektoren an. Dann ist

$$\begin{aligned}
 (a \cdot b) \cdot c &= \left(\begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} \right) \cdot \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} = \\
 &= \begin{pmatrix} a_1 \cdot b_1 - a_2 \cdot b_2 \\ a_1 \cdot b_2 + a_2 \cdot b_1 \end{pmatrix} \cdot \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} = \\
 &= \begin{pmatrix} (a_1 \cdot b_1 - a_2 \cdot b_2) \cdot c_1 - (a_1 \cdot b_2 + a_2 \cdot b_1) \cdot c_2 \\ (a_1 \cdot b_1 - a_2 \cdot b_2) \cdot c_2 + (a_1 \cdot b_2 + a_2 \cdot b_1) \cdot c_1 \end{pmatrix} = \\
 &= \begin{pmatrix} a_1 \cdot (b_1 \cdot c_1 - b_2 \cdot c_2) - a_2 \cdot (b_2 \cdot c_1 + b_1 \cdot c_2) \\ a_1 \cdot (b_1 \cdot c_2 - b_2 \cdot c_1) + a_2 \cdot (b_1 \cdot c_1 - b_2 \cdot c_2) \end{pmatrix} = \\
 &= \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot \begin{pmatrix} b_1 \cdot c_1 - b_2 \cdot c_2 \\ b_1 \cdot c_2 + b_2 \cdot c_1 \end{pmatrix} = \\
 &= \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot \left(\begin{pmatrix} b_1 \\ b_2 \end{pmatrix} \cdot \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} \right) = \\
 &= a \cdot (b \cdot c).
 \end{aligned}$$

Für den Beweis der Kommutativität der Multiplikation betrachten wir

$$\begin{aligned}
 (a_1, a_2) \cdot (b_1, b_2) &= (a_1 \cdot b_1 - a_2 \cdot b_2, a_1 \cdot b_2 + a_2 \cdot b_1) = \\
 &= (b_1 \cdot a_1 - b_2 \cdot a_2, b_1 \cdot a_2 + b_2 \cdot a_1) = \\
 &= (b_1, b_2) \cdot (a_1, a_2).
 \end{aligned}$$

Da $(a_1, a_2) \cdot (1, 0) = (a_1 \cdot 1 - a_2 \cdot 0, a_1 \cdot 0 + a_2 \cdot 1) = (a_1, a_2)$ ist, handelt es sich bei $(1, 0)$ um das Einselement.

Um die Existenz von Inversen für $(a_1, a_2) \in \mathbb{C}$ mit $(a_1, a_2) \neq (0, 0)$ zu beweisen, betrach-

ten wir die komplexe Zahl

$$a^{-1} := \left(\frac{a_1}{a_1^2 + a_2^2}, -\frac{a_2}{a_1^2 + a_2^2} \right).$$

Es ist

$$\begin{aligned} a \cdot a^{-1} &= (a_1, a_2) \cdot \left(\frac{a_1}{a_1^2 + a_2^2}, -\frac{a_2}{a_1^2 + a_2^2} \right) = \\ &= \left(a_1 \cdot \frac{a_1}{a_1^2 + a_2^2} + a_2 \cdot \frac{a_2}{a_1^2 + a_2^2}, -a_1 \cdot \frac{a_2}{a_1^2 + a_2^2} + a_2 \cdot \frac{a_1}{a_1^2 + a_2^2} \right) = \\ &= \left(\frac{a_1^2 + a_2^2}{a_1^2 + a_2^2}, \frac{a_1 \cdot a_2 - a_1 \cdot a_2}{a_1^2 + a_2^2} \right) = \\ &= (1, 0). \end{aligned}$$

Zuletzt beweisen wir die Distributivität:

$$\begin{aligned} a \cdot (b + c) &= (a_1, a_2) \cdot ((b_1, b_2) + (c_1, c_2)) = \\ &= (a_1, a_2) \cdot (b_1 + c_1, b_2 + c_2) = \\ &= (a_1 \cdot (b_1 + c_1) - a_2 \cdot (b_2 + c_2), a_1 \cdot (b_2 + c_2) + a_2 \cdot (b_1 + c_1)) = \\ &= (a_1 \cdot b_1 + a_1 \cdot c_1 - a_2 \cdot b_2 - a_2 \cdot c_2, a_1 \cdot b_2 + a_1 \cdot c_2 + a_2 \cdot b_1 + a_2 \cdot c_1) = \\ &= (a_1 \cdot b_1 - a_2 \cdot b_2, a_1 \cdot b_2 + a_2 \cdot b_1) + (a_1 \cdot c_1 - a_2 \cdot c_2, a_1 \cdot c_2 + a_2 \cdot c_1) = \\ &= (a_1, a_2) \cdot (b_1, b_2) + (a_1, a_2) \cdot (c_1, c_2) = \\ &= a \cdot b + a \cdot c. \end{aligned}$$

□

Für eine komplexe Zahl a schreiben wir das additive Inverse auch als $-a$ an, während wir das multiplikative Inverse von $a \neq (0,0)$ als a^{-1} notieren. Wie üblich vereinbaren wir, anstelle von $a + (-b)$ direkt $a - b$ zu schreiben.

Nachdem die komplexen Zahlen einen Körper bilden, folgen insbesondere die Eigenschaften aus Proposition 3.2.6.

Die komplexen Zahlen sind daher insgesamt eine zweidimensionale assoziative, kommutative, reelle Divisionsalgebra mit Einselement.

5.4 Einbettung von $\mathbb{R}$ in $\mathbb{C}$

Wir wollen nun beweisen, dass es sich bei den komplexen Zahlen tatsächlich um eine Erweiterung der reellen Zahlen handelt. Insbesondere wollen wir die gewohnte Schreibweise komplexer Zahlen zurückgewinnen.

Proposition 5.4.1. *Sei $\varphi : \mathbb{R} \rightarrow \mathbb{C}$ mit $\varphi(a) := (a, 0)$. Dann handelt es sich bei φ um eine Einbettung.*

Beweis. Seien $a, b \in \mathbb{R}$. Dann ist

$$\varphi(a + b) = (a + b, 0) = (a, 0) + (b, 0) = \varphi(a) + \varphi(b).$$

Zudem ist

$$\varphi(a \cdot b) = (a \cdot b, 0) = (a \cdot b - 0 \cdot 0, a \cdot 0 + 0 \cdot b) = (a, 0) \cdot (b, 0) = \varphi(a) \cdot \varphi(b).$$

Um die Injektivität von φ zu beweisen, nehmen wir $\varphi(a) = \varphi(b)$ an. Dann folgt durch komponentenweisen Vergleich direkt, dass $a = b$ ist. $\square$

Die reellen Zahlen lassen sich somit in die komplexen Zahlen einbetten. Um die gewohnte Schreibweise zu erhalten, vereinbaren wir, in der Notation einer komplexen Zahl (λ_1, λ_2) bezüglich der Standardbasis des $\mathbb{R}^2$ (siehe Abschnitt 5.2) anstelle von

$$\lambda_1 \cdot e + \lambda_2 \cdot i$$

direkt

$$\lambda_1 + \lambda_2 \cdot i$$

zu schreiben. Insbesondere erhalten wir damit, dass $i^2 = -1$ ist. Für $a, b \in \mathbb{C}$ mit $b \neq 0$ notieren wir $a \cdot b^{-1}$ als $\frac{a}{b}$.

Wir können nun zeigen, dass es keine Totalordnung auf den komplexen Zahlen gibt, die mit den Ordnungsaxiomen (Proposition 4.4.12) und den Rechenregeln für Ungleichungen (Proposition 4.4.13) auf den reellen Zahlen verträglich ist. Denn gäbe es eine

solche, so müsste jedenfalls $-1 < 0 < 1$ gelten. Wäre nun $i > 0$, so müsste aufgrund des zweiten Ordnungsaxioms $-1 = i^2 > 0$ sein, ein Widerspruch. Es muss also $i < 0$ sein, woraus nach den Rechenregeln für Ungleichungen $-i > 0$ folgt. Dann ist aufgrund des zweiten Ordnungsaxioms wieder $-1 = (-i)^2 > 0$, ein Widerspruch.

Für das Rechnen mit konjugiert komplexen Zahlen können wir folgende Resultate beweisen.

Proposition 5.4.2. Seien $a, b \in \mathbb{C}$. Dann ist $\overline{a + b} = \bar{a} + \bar{b}$ und $\overline{a \cdot b} = \bar{a} \cdot \bar{b}$. Weiters ist $a \cdot \bar{a} = |a|^2$, eine reelle Zahl.

Beweis. Wir schreiben $a = a_1 + i \cdot a_2$ und $b = b_1 + i \cdot b_2$. Dann ist

$$\begin{aligned}\overline{a + b} &= \overline{a_1 + b_1 + i \cdot (a_2 + b_2)} = \\ &= a_1 + b_1 - i \cdot (a_2 + b_2) = \\ &= (a_1 - i \cdot a_2) + (b_1 - i \cdot b_2) = \\ &= \bar{a} + \bar{b}.\end{aligned}$$

Zudem ist

$$\begin{aligned}\overline{a \cdot b} &= \overline{(a_1 \cdot b_1 - a_2 \cdot b_2) + i \cdot (a_1 \cdot b_2 + a_2 \cdot b_1)} = \\ &= (a_1 \cdot b_1 - a_2 \cdot b_2) - i \cdot (a_1 \cdot b_2 + a_2 \cdot b_1) = \\ &= (a_1 - i \cdot a_2) \cdot (b_1 - i \cdot b_2) = \\ &= \bar{a} \cdot \bar{b}.\end{aligned}$$

Für das Produkt $a \cdot \bar{a}$ erhalten wir

$$\begin{aligned}a \cdot \bar{a} &= (a_1 + i \cdot a_2) \cdot (a_1 - i \cdot a_2) = \\ &= a_1^2 + i \cdot a_1 \cdot a_2 - i \cdot a_1 \cdot a_2 + a_2^2 = \\ &= a_1^2 + a_2^2 = \\ &= |a|^2.\end{aligned}$$

□

5.5 Polardarstellung komplexer Zahlen

Betrachten wir die komplexe Zahl $2 + 3 \cdot i$ in Abbildung 5.1 genauer. Wir können diese offenbar durch Angabe des Winkels zur positiven x -Achse φ und des Abstands zum Nullpunkt r eindeutig angeben. Diese Darstellung einer komplexen Zahl durch ein Abstand-Winkel-Paar ist als Polardarstellung bekannt.

Definition 5.5.1 ([48, S. 328]). Sei $a + b \cdot i \neq 0$ eine komplexe Zahl. Dann nennen wir (r, φ) mit $r := |a + b \cdot i|$ und

$$\varphi := \begin{cases} \arctan \frac{b}{a} & \text{wenn } a > 0, \\ \frac{\pi}{2} & \text{wenn } a = 0 \text{ und } b > 0, \\ \arctan \frac{b}{a} + \pi & \text{wenn } a < 0, \\ -\frac{\pi}{2} & \text{wenn } a = 0 \text{ und } b < 0, \end{cases}$$

die Polardarstellung von $a + b \cdot i$.

Ist eine komplexe Zahl z in Polardarstellung (r, φ) gegeben, so erhalten wir die gewohnte Schreibweise über $z = r \cdot (\cos \varphi + i \cdot \sin \varphi)$.

Die Polardarstellung ist vor allem für das Multiplizieren komplexer Zahlen interessant. Haben wir $a, b \in \mathbb{C}$ mit $a = (r_1, \varphi_1)$ und $b = (r_2, \varphi_2)$, so können wir mit Hilfe der Winkelsummensätze das Produkt $a \cdot b$ berechnen:

$$\begin{aligned} a \cdot b &= r_1 \cdot (\cos \varphi_1 + i \cdot \sin \varphi_1) \cdot r_2 \cdot (\cos \varphi_2 + i \cdot \sin \varphi_2) = \\ &= r_1 \cdot r_2 \cdot (\cos(\varphi_1) \cdot \cos(\varphi_2) - \sin(\varphi_1) \cdot \sin(\varphi_2) + \\ &\quad + i \cdot (\cos(\varphi_1) \cdot \sin(\varphi_2) + \sin(\varphi_1) \cdot \cos(\varphi_2))) = \\ &= r_1 \cdot r_2 \cdot (\cos(\varphi_1 + \varphi_2) + i \cdot \sin(\varphi_1 + \varphi_2)). \end{aligned}$$

Das Produkt $a \cdot b$ entspricht somit in Polardarstellung der Zahl $(r_1 \cdot r_2, \varphi_1 + \varphi_2)$. Durch Induktion lässt sich zeigen, dass das Produkt $z_1 \cdot z_2 \cdot \dots \cdot z_n$ komplexer Zahlen z_i , welche jeweils die Polardarstellung (r_i, φ_i) besitzen, folgende Polardarstellung besitzt:

$$(r_1 \cdot r_2 \cdot \dots \cdot r_n, \varphi_1 + \varphi_2 + \dots + \varphi_n).$$

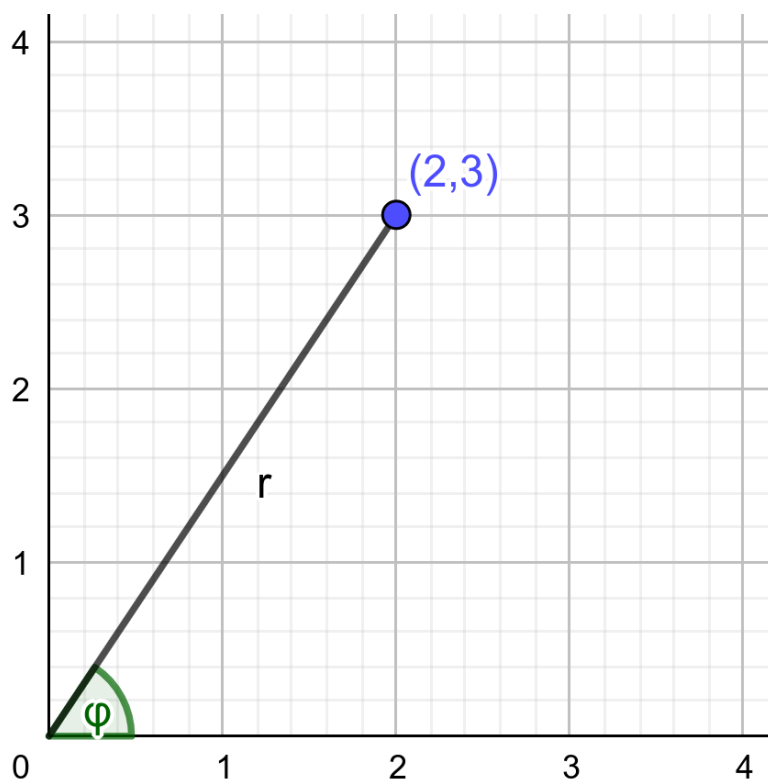


Abbildung 5.1: Die komplexe Zahl $2 + 3 \cdot i$ lässt über den Winkel φ und den Abstand r zum Nullpunkt beschreiben.

Sind daher $a, b \in \mathbb{C}$ komplexe Zahlen mit $|a| = 1$, so können wir die Multiplikation $a \cdot b$ geometrisch so verstehen, dass der Vektor b in der Gaußschen Zahlenebene um jenen Winkel, den a beschreibt, rotiert wird.

Wir können nun sehr einfach die Existenz komplexer Nullstellen der Gleichung $x^n = z$ beweisen.

Proposition 5.5.2. *Sei (r, φ) die Polardarstellung der komplexen Zahl z . Dann gibt es für $n \in \mathbb{N}$ mit $n \geq 1$ mindestens eine Lösung von $x^n = z$.*

Beweis. Wir setzen $\alpha := \sqrt[n]{r}$ und $x = \left(\alpha, \frac{\varphi}{n}\right)$. Durch Anwenden der zuvor gezeigten Rechenregel für die Multiplikation komplexer Zahlen in Polardarstellung erhalten wir $\left(\alpha, \frac{\varphi}{n}\right)^n = \left(\alpha^n, n \cdot \frac{\varphi}{n}\right) = (r, \varphi)$. $\square$

Wegen $\sin \varphi = \sin(\varphi + 2 \cdot \pi \cdot k)$ und $\cos \varphi = \cos(\varphi + 2 \cdot \pi \cdot k)$ für $k \in \mathbb{N}$, stellt auch $\left(\alpha, \frac{\varphi}{n} + \frac{2 \cdot \pi \cdot k}{n}\right)$ eine Nullstelle von $x^n = z$ dar, wobei es für $k \in \{0, 1, \dots, n-1\}$ genau n verschiedene Lösungen gibt. Wir können daher Proposition 5.5.2 wie folgt erweitern.

Proposition 5.5.3. Die komplexe Gleichung $x^n = z$ für $z \in \mathbb{C}$ und $n \in \mathbb{N}$ mit $n \geq 1$ besitzt genau n verschiedene komplexe Lösungen.

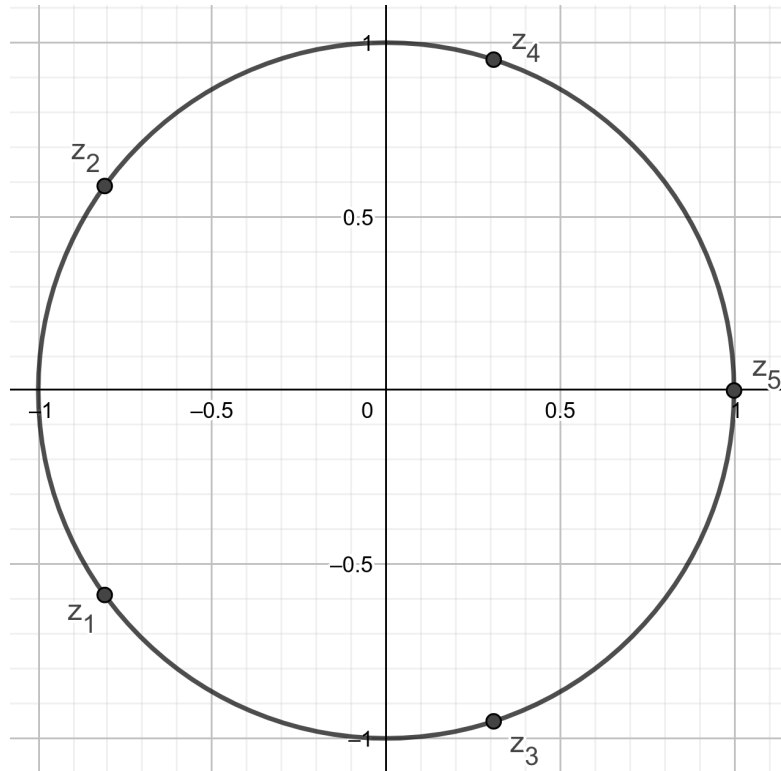


Abbildung 5.2: Darstellung der komplexen Lösungen der Gleichung $x^5 = 1$.

Zwischen der komplexen Exponentialfunktion und der Polardarstellung besteht folgender Zusammenhang.

Proposition 5.5.4. Für $\varphi \in \mathbb{R}$ ist $e^{i \cdot \varphi} = \cos \varphi + i \cdot \sin \varphi$.

Ist eine komplexe Zahl z in ihrer Polardarstellung (r, φ) gegeben, so gilt demnach $z = r \cdot e^{i \cdot \varphi}$. Speziell für die Zahl -1 mit der Polardarstellung $(1, \pi)$ ergibt sich $-1 = e^{i \cdot \pi}$, beziehungsweise $e^{i \cdot \pi} + 1 = 0$, die sogenannte eulersche Identität. Diese wird zu einer der

schönsten Formeln der Mathematik gezählt, da sie einen Zusammenhang zwischen der Eulerschen Zahl e , der imaginären Einheit i , der Kreiszahl π , sowie den grundlegenden Zahlen Null und Eins herstellt.

5.6 Der Fundamentalsatz der Algebra

Betrachten wir die quadratische Gleichung $x^2 + p \cdot x + q = 0$ mit $p, q \in \mathbb{R}$. Wie wir in Abschnitt 5.1 gesehen haben, können wir diese durch Addition von $\frac{p^2}{4} - q$ auf beiden Seiten auf ein vollständiges Quadrat ergänzen:

$$(x + p)^2 = \frac{p^2 - 4 \cdot q}{4}.$$

Im Fall $p^2 - 4 \cdot q > 0$ gibt es zwei reelle Lösungen, und für $p^2 - 4 \cdot q = 0$ eine doppelte reelle Lösung dieser Gleichung. Interessant ist nun $p^2 - 4 \cdot q < 0$: Während wir in den reellen Zahlen die Gleichung nicht lösen konnten, da Quadratzahlen stets größer als Null sind, gibt es nach Proposition 5.5.3 in den komplexen Zahlen genau zwei Lösungen dieser Gleichung. Es stellt sich somit die Frage, ob auch Gleichungen höherer Ordnung in den komplexen Zahlen stets lösbar sind. Die (positive) Antwort auf dieser Frage gibt der Fundamentalsatz der Algebra.

Zunächst benötigen wir folgende Proposition. Diese ermöglicht uns, ähnlich wie Proposition 2.5.3, eine Division mit Rest für Polynome.

Proposition 5.6.1. *Seien $p(z)$ und $q(z)$ komplexe Polynome mit $q(z) \neq 0$. Dann gibt es Polynome $s(z)$ und $r(z)$ mit*

$$p(z) = q(z) \cdot s(z) + r(z),$$

wobei entweder $r(z) = 0$ ist, oder $0 \leq \text{grad } r(z) < \text{grad } q(z)$ gilt.

Weiters benötigen wir folgende Resultate.

Proposition 5.6.2. *Sei (X, d) ein kompakter metrischer Raum, und $f : X \rightarrow \mathbb{R}$ eine stetige Funktion. Dann gibt es $x_1, x_2 \in X$ mit $f(x_1) \leq f(x) \leq f(x_2)$ für alle $x \in X$. Die Funktion nimmt daher auf X ein Minimum und ein Maximum an.*

Proposition 5.6.3. Für $r \in \mathbb{R}^+$ ist die Menge $\{z \in \mathbb{C} \mid |z| \leq r\}$ kompakt.

Das wesentliche Argument im Beweis des Fundamentalsatzes ist folgende Aussage, die auf den Mathematiker d'Alembert zurückgeht. Wir zeigen einen Beweis nach [1, S. 147–148].

Proposition 5.6.4 ([1, S. 147–148]). Sei $p(z) = a_n \cdot z^n + \dots + a_1 \cdot z + a_0$ ein nichtkonstantes komplexes Polynom, $p(z_0) \neq 0$, und $R \in \mathbb{R}^+$. Dann gibt es für $C = \{z \in \mathbb{C} \mid |z - z_0| < R\}$ ein $z_1 \in C$ mit $|p(z_1)| < |p(z_0)|$.

Beweis. Wir zeigen zunächst, dass es ein $c \in \mathbb{C}$ mit $c \neq 0$ und ein Polynom $r(w)$ mit $r(0) = 0$ gibt, sodass wir

$$p(z_0 + w) = p(z_0) + c \cdot w^m \cdot (1 + r(w))$$

schreiben können. Aufgrund des binomischen Lehrsatzes und den Regeln für das Vertauschen von Summen gilt:

$$\begin{aligned} p(z_0 + w) &= \sum_{k=0}^n a_k \cdot (z_0 + w)^k = \\ &= \sum_{k=0}^n a_k \cdot \sum_{l=0}^k \binom{k}{l} \cdot w^l \cdot z_0^{k-l} = \\ &= \sum_{k=0}^n \cdot \sum_{l=0}^k a_k \cdot \binom{k}{l} \cdot w^l \cdot z_0^{k-l} = \\ &= \sum_{l=0}^n \cdot \sum_{k=l}^n a_k \cdot \binom{k}{l} \cdot w^l \cdot z_0^{k-l} = \\ &= \sum_{l=0}^n w^l \cdot \sum_{k=l}^n \binom{k}{l} \cdot a_k \cdot z_0^{k-l} = \\ &= p(z_0) + \underbrace{\sum_{l=1}^n w^l \cdot \sum_{k=l}^n \binom{k}{l} \cdot a_k \cdot z_0^{k-l}}_{=: q(w)}. \end{aligned}$$

Wir können nun aus $q(w)$ die kleinste Potenz w^m mit einem Koeffizienten $c \neq 0$ herausheben, und erhalten dadurch $q(w) = c \cdot w^m \cdot (1 + r(w))$ für ein Polynom $r(w)$ mit $r(0) = 0$.

Es gilt also $p(z_0 + w) = p(z_0) + c \cdot w^m \cdot (1 + r(w))$. Ist nun

$$\eta := \sqrt[m]{\left| \frac{p(z_0)}{c} \right|},$$

so ist für $|w| < \eta$ sicher

$$|c \cdot w^m| = |c| \cdot |w|^m < |c| \cdot \eta^m = |c \cdot \eta^m| = |p(z_0)|. \quad (5.4)$$

Da es sich bei $r(w)$ um eine Polynomfunktion handelt, ist diese stetig. Es gibt daher ein $\delta > 0$, sodass für $|w - 0| < \delta$ stets $|r(w) - r(0)| < 1$ ist. Wegen $r(0) = 0$ folgt daher aus $|w| < \delta$, dass $|r(w)| < 1$ ist.

Insgesamt gilt also: Ist $S := \min\{\delta, \eta\}$, so ist für $|w| < S$ stets $|c \cdot w^m| < |p(z_0)|$ und $|r(w)| < 1$.

Wegen Proposition 5.5.2 gibt es ein $\zeta \in \mathbb{C}$ mit

$$\zeta^m = -\frac{\frac{p(z_0)}{c}}{\left| \frac{p(z_0)}{c} \right|}.$$

Durch Anwenden der Rechenregeln für Beträge erhalten wir $|\zeta^m| = 1$, und damit auch $|\zeta| = 1$. Sei außerdem $0 < \varepsilon < \min\{R, S\}$. Wir behaupten nun, dass $z_1 := z_0 + \varepsilon \cdot \zeta$ die gesuchte Zahl mit $|p(z_1)| < |p(z_0)|$ ist. Wegen

$$|z_1 - z_0| = |\varepsilon \cdot \zeta| = \underbrace{|\varepsilon|}_{< R} \cdot \underbrace{|\zeta|}_{=1} < R$$

ist jedenfalls $z_1 \in C$. Weiters ist für ein geeignetes $t \in \mathbb{R}^+$

$$c \cdot \varepsilon^m \cdot \zeta^m = -c \cdot \varepsilon^m \cdot \frac{\frac{p(z_0)}{c}}{\left| \frac{p(z_0)}{c} \right|} = -\frac{\varepsilon^m \cdot |c|}{|p(z_0)|} \cdot p(z_0) = -t \cdot p(z_0),$$

wobei wegen $\varepsilon < S \leq \eta$ aus Gleichung 5.4 folgt, dass $t < 1$ ist. Wir können daher

$|p(z_1)| = |p(z_0 + \varepsilon \cdot \zeta)|$ wie folgt abschätzen:

$$\begin{aligned}
 |p(z_0 + \varepsilon \cdot \zeta)| &= |p(z_0) + c \cdot (\varepsilon \cdot \zeta)^m \cdot (1 + r(\varepsilon \cdot \zeta))| = \\
 &= |p(z_0) + \underbrace{c \cdot \varepsilon^m \cdot \zeta^m}_{=-t \cdot p(z_0)} \cdot (1 + r(\varepsilon \cdot \zeta))| = \\
 &= |p(z_0) - t \cdot p(z_0) \cdot (1 + r(\varepsilon \cdot \zeta))| = \\
 &= |p(z_0) \cdot (1 - t) - t \cdot p(z_0) \cdot r(\varepsilon \cdot \zeta)| \leq \\
 &\leq |p(z_0)| \cdot (1 - t) + t \cdot |p(z_0)| \cdot \underbrace{|r(\varepsilon \cdot \zeta)|}_{<1} < \\
 &< |p(z_0)| \cdot (1 - t) + t \cdot |p(z_0)| = \\
 &= |p(z_0)|.
 \end{aligned}$$

Somit ist $|p(z_1)| < |p(z_0)|$. □

Satz 5.6.5 ([1, S. 147–149; 12, S. 25–26]). Sei $p(z) = a_n \cdot z^n + a_{n-1} \cdot z^{n-1} + \dots + a_0$ ein nichtkonstantes komplexes Polynom. Dann gibt es ein $z_0 \in \mathbb{C}$ mit $p(z_0) = 0$.

Beweis. Für den Beweis zeigen wir, dass die reellwertige Funktion $|p(z)|$ ein globales Minimum annimmt, und dieses gleich Null ist.

Im ersten Schritt beweisen wir, dass wir eine Kreisscheibe $C = \{z \in \mathbb{C} \mid |z| \leq R\}$ mit Radius R um den Nullpunkt finden können, außerhalb derer $|p(z)|$ stets größer als $1 + |p(0)|$ ist. Für $z \neq 0$ ist

$$\begin{aligned}
 |p(z)| &= |a_n \cdot z^n + a_{n-1} \cdot z^{n-1} + \dots + a_0| = \\
 &= |a_n \cdot z^n| \cdot \left| 1 + \frac{a_{n-1}}{a_n} \cdot \frac{1}{z} + \frac{a_{n-2}}{a_n} \cdot \frac{1}{z^2} + \dots + \frac{a_0}{a_n} \cdot \frac{1}{z^n} \right|.
 \end{aligned}$$

Für z gegen Unendlich strebt der erste Faktor gegen Plus Unendlich, und der zweite gegen Eins. Daher strebt $|p(z)|$ ebenfalls gegen Plus Unendlich, und es gibt ein $R \in \mathbb{R}^+$ mit $|p(z)| > 1 + |p(0)|$ für alle $z \in \mathbb{C}$ mit $|z| > R$.

Definieren wir nun die Kreisscheibe $C := \{z \in \mathbb{C} \mid |z| \leq R\}$. Nach Proposition 5.6.3 ist diese kompakt, und da es sich bei $|p(z)|$ um eine stetige Funktion handelt, gibt es nach Proposition 5.6.2 ein $z_0 \in C$ mit $|p(z_0)| \leq |p(z)|$ für alle $z \in C$. Wir haben daher ein

Minimum auf der Kreisscheibe gefunden. Für $z \in \mathbb{C}$ mit $z \notin C$ ist $|z| > R$, und daher $|p(z)| > 1 + |p(0)|$. Da $0 \in C$ ist, gilt aber $|p(z_0)| \leq |p(0)| < 1 + |p(0)| < |p(z)|$. Die Funktion $|p(z)|$ nimmt daher in z_0 ein globales Minimum an.

Ist $p(z_0) = 0$, so sind wir fertig. Andernfalls gibt es nach Proposition 5.6.4 ein $z_1 \in C$ mit $|p(z_1)| < |p(z_0)|$, im Widerspruch zur Minimalität von $|p(z_0)|$. Es muss somit $p(z_0) = 0$ sein. $\square$

Vom Fundamentalsatz der Algebra sind zahlreiche Beweise bekannt. Besonders elegante und kurze Beweise gelingen über Sätze aus der komplexen Analysis, wie etwa über den Satz von Liouville.

Im folgenden Abschnitt zeigen wir einige Resultate, die sich aus dem Fundamentalsatz der Algebra folgern lassen.

Ist $p(z)$ ein nichtkonstantes Polynom mit einer Nullstelle z_0 , so können wir diese vom Polynom abspalten.

Proposition 5.6.6. *Sei $p(z)$ ein nichtkonstantes komplexes Polynom, und $z_0 \in \mathbb{C}$ mit $p(z_0) = 0$. Dann gibt es ein Polynom $q(z)$ mit $\text{grad } q(z) = \text{grad } p(z) - 1$, sodass*

$$p(z) = (z - z_0) \cdot q(z).$$

Beweis. Da wir eine Polynomdivision durchführen können, gibt es komplexe Polynome $q(z)$ und $r(z)$ mit $r(z) = 0$ oder $\text{grad } r(z) = 0$, sodass $p(z) = q(z) \cdot (z - z_0) + r(z)$ ist. Wegen $\text{grad } r(z) = 0$ oder $r(z) = 0$, muss $r(z)$ auf ganz $\mathbb{C}$ konstant sein. Es ist

$$0 = p(z_0) = (z_0 - z_0) \cdot p(z_0) + r(z_0) = r(z_0),$$

und daher $r(z) = 0$ für alle $z \in \mathbb{C}$. Damit ist

$$p(z) = (z - z_0) \cdot q(z).$$

Weiters ist

$$\text{grad } p(z) = \text{grad } ((z - z_0) \cdot q(z)) = \text{grad}(z - z_0) + \text{grad } q(z) = 1 + \text{grad } q(z),$$

womit $\text{grad } q(z) = \text{grad } p(z) - 1$ gezeigt ist. $\square$

Mit Hilfe der vorigen Proposition lässt sich nun zeigen, dass jedes nichtkonstante komplexe Polynom über den komplexen Zahlen in Linearfaktoren zerfällt.

Proposition 5.6.7. *Sei $p(z)$ ein nichtkonstantes komplexes Polynom vom Grad n . Dann gibt es n Linearfaktoren $z - z_i$ und ein $c \in \mathbb{C}$, mit*

$$p(z) = c \cdot \prod_{k=1}^n (z - z_k).$$

Zählt man mehrfache Nullstellen mit, besitzt das Polynom also genau n Nullstellen.

Beweis. Nach dem Fundamentalsatz der Algebra besitzt $p(z)$ eine Nullstelle z_1 . Diese können wir nach Proposition 5.6.6 abspalten, und erhalten $p(z) = (z - z_1) \cdot q_1(z)$ für ein Polynom $q_1(z)$ vom Grad $n - 1$.

Ist $\text{grad } q_1(z) = n - 1 = 0$, so sind wir fertig. Andernfalls gibt es eine Nullstelle z_2 von $q_1(z)$, die wegen $p(z) = (z - z_1) \cdot q_1(z)$ auch eine Nullstelle von $p(z)$ ist. Abspalten dieser Nullstelle liefert $q_1(z) = (z - z_2) \cdot q_2(z)$ und $p(z) = (z - z_1) \cdot (z - z_2) \cdot q_2(z)$, wobei $q_2(z)$ vom Grad $n - 2$ ist.

Durch Wiederholen der vorigen Schritte erhalten wir eine Kette von n Linearfaktoren mit

$$p(z) = (z - z_1) \cdot (z - z_2) \cdot \dots \cdot (z - z_n) \cdot \underbrace{q_n(z)}_{=: c},$$

wobei $q_n(z)$ vom Grad $n - n = 0$, also konstant ist. $\square$

Angenommen, $p(z)$ wäre ein reelles Polynom vom Grad n mit der (komplexen) Nullstelle z_0 . Wegen der Rechenregeln für konjugiert komplexe Zahlen gilt

$$\begin{aligned} p(\overline{z_0}) &= \sum_{k=0}^n a_k \cdot (\overline{z_0})^k = \\ &= \sum_{k=0}^n a_k \cdot \overline{z_0^k} = \\ &= \sum_{k=0}^n \overline{a_k \cdot z_0^k} = \end{aligned}$$

$$\begin{aligned} &= \overline{\sum_{k=0}^n a_n \cdot z_0^k} = \\ &= \bar{0} = 0. \end{aligned}$$

Es handelt sich also auch bei $\bar{z}_0$ um eine Nullstelle von $p(z)$, womit $z - z_0$ und $z - \bar{z}_0$ Linearfaktoren von $p(z)$ sind. Multiplizieren dieser Linearfaktoren liefert

$$(z - z_0) \cdot (z - \bar{z}_0) = z^2 - z \cdot \bar{z}_0 - z \cdot z_0 + z_0 \cdot \bar{z}_0 = z^2 - z \cdot \underbrace{(z_0 + \bar{z}_0)}_{\in \mathbb{R}} + \underbrace{z_0 \cdot \bar{z}_0}_{\in \mathbb{R}}.$$

Wir erhalten daher ein reelles quadratisches Polynom, das über den reellen Zahlen keine Nullstelle besitzt. Da für jede komplexe Nullstelle z_0 von $p(z)$ auch die konjugiert komplexe Zahl $\bar{z}_0$ eine Nullstelle darstellt, kommen in der Linearfaktorzerlegung von $p(z)$ ausschließlich Produkte wie das obige vor, die sich auf quadratische reelle Polynome erweitern lassen. Es gilt damit folgende Proposition.

Proposition 5.6.8 ([12, S. 27]). *Sei $p(z)$ ein reelles Polynom. Dann lässt sich $p(z)$ als Produkt von linearen und quadratischen reellen Polynomen anschreiben. Die quadratischen Polynome besitzen dabei keine Nullstellen über den reellen Zahlen.*

6 Die Quaternionen

Ausgangspunkt für die Entdeckung der Quaternionen sind Überlegungen von William Hamilton zur geometrischen Veranschaulichung von Addition und Multiplikation im Raum. Wir haben gesehen, dass sich mit komplexen Zahlen Koordinaten in der Ebene beschreiben lassen, und Rechenoperationen auf diesen eine geometrische Deutung besitzen (etwa Rotationen, siehe Abschnitt 5.5). Hamilton stellte sich nun die Frage, ob man nicht eine Erweiterung der komplexen Zahlen ins Dreidimensionale finden könne, mit deren Hilfe sich geometrische Operationen im Raum einfach beschreiben lassen [20, S. 155–156].

Hamilton suchte jahrelang vergeblich nach einer solchen Erweiterung. Heute lässt sich die Ergebnislosigkeit seiner Suche damit erklären, dass es schlichtweg keine dreidimensionale Divisionsalgebra über den reellen Zahlen gibt. In Proposition 6.4.1 geben wir einen allgemeinen Beweis dafür, dass eine reelle Divisionsalgebra mit ungerader Dimension nur im Fall der reellen Zahlen möglich ist.

Während eines Spaziergangs am 16. Oktober 1843 kam Hamilton nun auf die Idee, nach einer Erweiterung im Vierdimensionalen zu suchen, und ritzte die für die Multiplikation grundlegenden Gleichungen in den Stein der Brougham Bridge ein [20, S. 156–157].

Hamilton war zeit seines Lebens davon überzeugt, dass die *Quaternionen*, wie er diesen Zahlenbereich nannte, von ähnlicher Bedeutung für die Mathematik wären wie die Infinitesimalrechnung. Tatsächlich hat er die innermathematische Bedeutung der Quaternionen stark überschätzt. Heute wird von ihnen vor allem in mathematischen Anwendungsfeldern Gebrauch gemacht. So lassen sich etwa in der Physik vierdimensionale Raum-Zeit-Gleichungen durch Quaternionen einfacher formulieren. Auch in der Computergrafik werden Quaternionen verwendet, da sich durch sie Rotationen im Raum effizient beschreiben lassen [20, S. 157–158; 15, S. xii].

6.1 Konstruktion der Quaternionen

Hamilton verstand die Quaternionen als vierdimensionale reelle Zahlen. Wir könnten daher prinzipiell die Quaternionen als $\mathbb{R}^4$ definieren, und auf diesem Vektorraum eine Multiplikation einführen. Die Darstellung als vierdimensionaler Vektor reeller Zahlen ist in der Praxis aber eher mühsam: Die Anzahl der Variablen wird schnell sehr groß und Beweise dadurch schwer lesbar und nachvollziehbar. Wir können allerdings $\mathbb{R}^4$ als $\mathbb{R}^2 \times \mathbb{R}^2$ verstehen. Da wir die komplexen Zahlen als $\mathbb{R}^2$ definiert haben, wäre eine Definition der Quaternionen als $\mathbb{C} \times \mathbb{C}$ überlegenswert.

Um diese Definition zu motivieren, gehen wir zunächst (weniger formell) von Hamiltons Definition der Quaternionen aus. Ähnlich wie sich komplexe Zahlen als Summe von Real- und Imaginärteil anschreiben lassen, hat Hamilton die Quaternionen als Zahlen der Form $a + b \cdot i + c \cdot j + d \cdot k$ definiert. Dabei sind a, b, c, d reelle Zahlen, i die bekannte imaginäre Einheit, und j sowie k zwei neue weitere imaginäre Einheiten. Die grundlegenden Beziehungen zwischen diesen Einheiten hat Hamilton in den Stein der Brougham Bridge eingeritzt; wir geben sie hier als Tabelle an.

$\cdot$	1	i	j	k
1	1	i	j	k
i	i	-1	k	$-j$
j	j	$-k$	-1	i
k	k	j	$-i$	-1

Tabelle 6.1: Multiplikationstabelle der Quaternionen [20, S. 159].

Die Multiplikationstabelle lässt sich grafisch sehr einfach darstellen, siehe Abbildung 6.1: Multipliziert man zwei Einheiten im Uhrzeigersinn, erhält man die den beiden nachfolgende Einheit; im Gegenuhrzeigersinn ist es das additive Inverse.

Betrachten wir nun ein Quaternion $a + b \cdot i + c \cdot j + d \cdot k$, und nehmen an, die Distributivität und Assoziativität der Multiplikation wären erfüllt. Aus dem Term $c \cdot j + d \cdot k$ können wir j herausheben, und erhalten mit Hilfe der Multiplikationstabelle $(c + d \cdot i) \cdot j$. Das Quaternion lässt sich daher als

$$(a + b \cdot i) + (c + d \cdot i) \cdot j$$

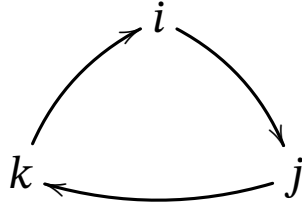


Abbildung 6.1: Grafische Darstellung der Multiplikation in $\mathbb{H}$ [2, S. 151].

anschreiben, wobei $a + b \cdot i$ und $c + d \cdot i$ komplexe Zahlen sind. Die Quaternionen lassen sich damit tatsächlich als Paare komplexer Zahlen verstehen. Dadurch wird folgende Definition gerechtfertigt.

Definition 6.1.1 ([20, S. 159]). Die Menge der Quaternionen $\mathbb{H}$ ist definiert als

$$\mathbb{H} := \mathbb{C} \times \mathbb{C} = \{(a, b) \mid a, b \in \mathbb{C}\}.$$

Wir setzen $e = (1, 0)$, $i = (i, 0)$, $j = (0, 1)$ und $k = (0, i)$, um die von Hamilton eingeführten Einheiten zu erhalten. Fassen wir die Quaternionen als $\mathbb{R}$ -Vektorraum auf, so bilden die Elemente $\{e, i, j, k\}$ offenbar eine Basis. Wir können also jedes Quaternion auch in der Form

$$\lambda_1 \cdot e + \lambda_2 \cdot i + \lambda_3 \cdot j + \lambda_4 \cdot k$$

mit $\lambda_i \in \mathbb{R}$ für $1 \leq i \leq 4$ anschreiben.

Analog zu den komplexen Zahlen definieren wir für ein Quaternion die zugehörige konjugierte Zahl und dessen Betrag.

Definition 6.1.2 ([2, S. 154]). Für $(a, b) \in \mathbb{H}$ definieren wir

$$\overline{(a, b)} := (\bar{a}, -b).$$

Definition 6.1.3 ([20, S. 169]). Für $(a, b) \in \mathbb{H}$ definieren wir

$$|(a, b)| := \sqrt{|a|^2 + |b|^2}.$$

Es lässt sich wieder beweisen, dass die Betragsfunktion alle bekannten Eigenschaften aufweist; unter anderem ist die Produktregel erfüllt [20, S. 171].

6.2 Rechenoperationen auf $\mathbb{H}$

Wie wir in diesem Abschnitt sehen werden, müssen wir für die Erweiterung auf die Quaternionen erstmals eine grundlegende Struktureigenschaft aufgeben: Die Multiplikation ist in $\mathbb{H}$ nicht kommutativ.

Zunächst zur Definition der Addition. Wie schon in den komplexen Zahlen wird diese komponentenweise durchgeführt.

Definition 6.2.1 ([20, S. 211]). Die Addition $+$: $\mathbb{H} \times \mathbb{H} \rightarrow \mathbb{H}$ ist definiert durch

$$(a, b) + (c, d) := (a + c, b + d).$$

Definition 6.2.2 ([20, S. 211]). Die Multiplikation $\cdot$: $\mathbb{H} \times \mathbb{H} \rightarrow \mathbb{H}$ ist definiert durch

$$(a, b) \cdot (c, d) := (a \cdot c - b \cdot \bar{d}, a \cdot d + b \cdot \bar{c}).$$

Vergleichen wir diese Definition mit der Multiplikation in den komplexen Zahlen, so fällt auf, dass die Definitionen nahezu identisch sind. Da eine reelle Zahl stets gleich ihrer konjugierten Zahl ist, könnten wir diese Definition auch für die Multiplikation in den komplexen Zahlen verwenden. Dieses Verfahren, das auch als *Cayley-Dickson-Konstruktion* bekannt ist, wird uns im nächsten Kapitel auf die Oktonionen führen [2, S. 154].

Durch Anwenden der Definition auf die Einheiten e, i, j, k lässt sich leicht verifizieren, dass die soeben definierte Multiplikation mit Tabelle 6.1 kompatibel ist.

Wie bereits am Anfang dieses Abschnitts erwähnt, ist die Multiplikation in den Quaternionen nicht mehr kommutativ. Das war bereits aus Tabelle 6.1 ersichtlich, und wird durch die in der Definition vorkommenden konjugiert komplexen Zahlen nochmals deutlich. Tatsächlich sind aber alle anderen Körperaxiome erfüllt. Eine algebraische

Struktur, welche alle Körperaxiome bis auf die Kommutativität der Multiplikation erfüllt, wird Schiefkörper genannt [34, S. 31].

Proposition 6.2.3. $(\mathbb{H}, +, \cdot, (0,0), (1,0))$ ist ein Schiefkörper. Es sind daher folgende Eigenschaften erfüllt.

- (1) $\forall a, b, c \in \mathbb{H} : (a + b) + c = a + (b + c)$ (Assoziativität von $+$).
- (2) $\forall a, b \in \mathbb{H} : a + b = b + a$ (Kommutativität von $+$).
- (3) $\forall a \in \mathbb{H} : a + (0,0) = a$ (Nullelement).
- (4) $\forall a \in \mathbb{H} : \exists -a \in \mathbb{H} : a + (-a) = (0,0)$ (Invertierbarkeit von $+$).
- (5) $\forall a, b, c \in \mathbb{H} : (a \cdot b) \cdot c = a \cdot (b \cdot c)$ (Assoziativität von $\cdot$).
- (6) $\forall a \in \mathbb{H} : a \cdot (1,0) = (1,0) \cdot a = a$ (Einselement).
- (7) $\forall a \in \mathbb{H} \setminus \{(0,0)\} : \exists a^{-1} \in \mathbb{H} : a \cdot a^{-1} = a^{-1} \cdot a = (1,0)$ (Invertierbarkeit von $\cdot$).
- (8) $\forall a, b, c \in \mathbb{H} : a \cdot (b + c) = a \cdot b + a \cdot c$ und $(b + c) \cdot a = b \cdot a + c \cdot a$ (Distributivität).

Beweis. Seien $a = (a_1, a_2)$, $b = (b_1, b_2)$ und $c = (c_1, c_2)$ Quaternionen. Die Eigenschaften (1) bis (4) folgen direkt daraus, dass die Addition komponentenweise definiert ist, und es sich bei den komplexen Zahlen um einen Körper handelt.

Zur Assoziativität der Multiplikation: Zwecks besserer Übersicht schreiben wir Quaternionen als Spaltenvektoren an. Dann ist

$$\begin{aligned}
 (a \cdot b) \cdot c &= \left(\begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} \right) \cdot \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} = \\
 &= \begin{pmatrix} a_1 \cdot b_1 - a_2 \cdot \overline{b_2} \\ a_1 \cdot b_2 + a_2 \cdot \overline{b_1} \end{pmatrix} \cdot \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} = \\
 &= \begin{pmatrix} (a_1 \cdot b_1 - a_2 \cdot \overline{b_2}) \cdot c_1 - (a_1 \cdot b_2 + a_2 \cdot \overline{b_1}) \cdot \overline{c_2} \\ (a_1 \cdot b_1 - a_2 \cdot \overline{b_2}) \cdot c_2 + (a_1 \cdot b_2 + a_2 \cdot \overline{b_1}) \cdot \overline{c_1} \end{pmatrix} = \\
 &= \begin{pmatrix} a_1 \cdot b_1 \cdot c_1 - a_2 \cdot \overline{b_2} \cdot c_1 - a_1 \cdot b_2 \cdot \overline{c_2} - a_2 \cdot \overline{b_1} \cdot \overline{c_2} \\ a_1 \cdot b_1 \cdot c_2 - a_2 \cdot \overline{b_2} \cdot c_2 + a_1 \cdot b_2 \cdot \overline{c_1} + a_2 \cdot \overline{b_1} \cdot \overline{c_1} \end{pmatrix} =
 \end{aligned}$$

$$\begin{aligned}
&= \begin{pmatrix} a_1 \cdot (b_1 \cdot c_1 - b_2 \cdot \overline{c_2}) - a_2 \cdot (\overline{b_1} \cdot \overline{c_2} + \overline{b_2} \cdot c_1) \\ a_1 \cdot (b_1 \cdot c_2 + b_2 \cdot \overline{c_1}) + a_2 \cdot (\overline{b_1} \cdot \overline{c_1} - \overline{b_2} \cdot c_2) \end{pmatrix} = \\
&= \begin{pmatrix} a_1 \cdot (b_1 \cdot c_1 - b_2 \cdot \overline{c_2}) - a_2 \cdot \overline{(b_1 \cdot c_2 + b_2 \cdot \overline{c_1})} \\ a_1 \cdot (b_1 \cdot c_2 + b_2 \cdot \overline{c_1}) + a_2 \cdot \overline{(b_1 \cdot c_1 - b_2 \cdot \overline{c_2})} \end{pmatrix} = \\
&= \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot \begin{pmatrix} b_1 \cdot c_1 - b_2 \cdot \overline{c_2} \\ b_1 \cdot c_2 + b_2 \cdot \overline{c_1} \end{pmatrix} = \\
&= \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot \left(\begin{pmatrix} b_1 \\ b_2 \end{pmatrix} \cdot \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} \right) = \\
&= a \cdot (b \cdot c).
\end{aligned}$$

Durch Einsetzen in die Definition der Multiplikation sehen wir direkt, dass es sich bei $(1,0)$ um das Einselement der Multiplikation handelt:

$$\begin{aligned}
(a_1, a_2) \cdot (1, 0) &= (a_1 \cdot 1 + a_2 \cdot 0, a_1 \cdot 0 + a_2 \cdot 1) = (a_1, a_2), \\
(1, 0) \cdot (a_1, a_2) &= (1 \cdot a_1 + 0 \cdot \overline{a_2}, 1 \cdot a_2 + 0 \cdot \overline{a_1}) = (a_1, a_2).
\end{aligned}$$

Um die Invertierbarkeit der Multiplikation zu zeigen, nehmen wir an, dass $(a_1, a_2) \in \mathbb{H}$ mit $(a_1, a_2) \neq (0, 0)$ ist. Wir setzen, ähnlich wie im entsprechenden Beweis für die komplexen Zahlen,

$$a^{-1} = \left(\frac{\overline{a_1}}{|a_1|^2 + |a_2|^2}, \frac{-a_2}{|a_1|^2 + |a_2|^2} \right).$$

Dann ist

$$\begin{aligned}
a \cdot a^{-1} &= (a_1, a_2) \cdot \left(\frac{\overline{a_1}}{|a_1|^2 + |a_2|^2}, \frac{-a_2}{|a_1|^2 + |a_2|^2} \right) = \\
&= \left(a_1 \cdot \frac{\overline{a_1}}{|a_1|^2 + |a_2|^2} - a_2 \cdot \frac{-a_2}{|a_1|^2 + |a_2|^2}, a_1 \cdot \frac{-a_2}{|a_1|^2 + |a_2|^2} + a_2 \cdot \frac{\overline{a_1}}{|a_1|^2 + |a_2|^2} \right) = \\
&= \left(\frac{|a_1|^2}{|a_1|^2 + |a_2|^2} - a_2 \cdot \frac{-\overline{a_2}}{|a_1|^2 + |a_2|^2}, \frac{-a_1 \cdot a_2}{|a_1|^2 + |a_2|^2} + \frac{a_2 \cdot a_1}{|a_1|^2 + |a_2|^2} \right) = \\
&= \left(\frac{|a_1|^2}{|a_1|^2 + |a_2|^2} + \frac{|a_2|^2}{|a_1|^2 + |a_2|^2}, 0 \right) = \\
&= (1, 0),
\end{aligned}$$

wobei wir für den Beweis verwendet haben, dass für $a \in \mathbb{C}$ stets $a \cdot \bar{a} = |a|^2$ ist. Um zu zeigen, dass es sich dabei auch um das Linksinverse handelt, betrachten wir

$$\begin{aligned}
 a^{-1} \cdot a &= \left(\frac{\bar{a}_1}{|a_1|^2 + |a_2|^2}, \frac{-a_2}{|a_1|^2 + |a_2|^2} \right) \cdot (a_1, a_2) = \\
 &= \left(\frac{\bar{a}_1}{|a_1|^2 + |a_2|^2} \cdot a_1 - \frac{-a_2}{|a_1|^2 + |a_2|^2} \cdot \bar{a}_2, \frac{\bar{a}_1}{|a_1|^2 + |a_2|^2} \cdot a_2 + \frac{-a_2}{|a_1|^2 + |a_2|^2} \cdot \bar{a}_1 \right) = \\
 &= \left(\frac{|a_1|^2}{|a_1|^2 + |a_2|^2} + \frac{|a_2|^2}{|a_1|^2 + |a_2|^2}, \frac{\bar{a}_1 \cdot a_2}{|a_1|^2 + |a_2|^2} - \frac{a_2 \cdot \bar{a}_1}{|a_1|^2 + |a_2|^2} \right) = \\
 &= (1, 0).
 \end{aligned}$$

Zuletzt zeigen wir die Gültigkeit der Distributivgesetze. Es ist

$$\begin{aligned}
 (a_1, a_2) \cdot ((b_1, b_2) + (c_1, c_2)) &= \\
 &= (a_1, a_2) \cdot (b_1 + c_1, b_2 + c_2) = \\
 &= (a_1 \cdot (b_1 + c_1) - a_2 \cdot \overline{(b_2 + c_2)}, a_1 \cdot (b_2 + c_2) + a_2 \cdot \overline{(b_1 + c_1)}) = \\
 &= (a_1 \cdot b_1 + a_1 \cdot c_1 - a_2 \cdot \bar{b}_2 - a_2 \cdot \bar{c}_2, a_1 \cdot b_2 + a_1 \cdot c_2 + a_2 \cdot \bar{b}_1 + a_2 \cdot \bar{c}_1) = \\
 &= (a_1 \cdot b_1 - a_2 \cdot \bar{b}_2, a_1 \cdot b_2 + a_2 \cdot \bar{b}_1) + (a_1 \cdot c_1 - a_2 \cdot \bar{c}_2, a_1 \cdot c_2 + a_2 \cdot \bar{c}_1) = \\
 &= (a_1, a_2) \cdot (b_1, b_2) + (a_1, a_2) \cdot (c_1, c_2),
 \end{aligned}$$

wobei wir für den Beweis die Distributivität der komplexen Zahlen verwendet haben. Das zweite Distributivgesetz lässt sich analog beweisen.

$$\begin{aligned}
 ((b_1, b_2) + (c_1, c_2)) \cdot (a_1, a_2) &= \\
 &= (b_1 + c_1, b_2 + c_2) \cdot (a_1, a_2) = \\
 &= ((b_1 + c_1) \cdot a_1 - (b_2 + c_2) \cdot \bar{a}_2, (b_1 + c_1) \cdot a_2 + (b_2 + c_2) \cdot \bar{a}_1) = \\
 &= (b_1 \cdot a_1 + c_1 \cdot a_1 - b_2 \cdot \bar{a}_2 - c_2 \cdot \bar{a}_2, b_1 \cdot a_2 + c_1 \cdot a_2 + b_2 \cdot \bar{a}_1 + c_2 \cdot \bar{a}_1) = \\
 &= (b_1 \cdot a_1 - b_2 \cdot \bar{a}_2, b_1 \cdot a_2 + b_2 \cdot \bar{a}_1) + (c_1 \cdot a_1 - c_2 \cdot \bar{a}_2, c_1 \cdot a_2 + c_2 \cdot \bar{a}_1) = \\
 &= (b_1, b_2) \cdot (a_1, a_2) + (c_1, c_2) \cdot (a_1, a_2). \quad \square
 \end{aligned}$$

Wie üblich schreiben wir $-a$ für das additive Inverse von a , und $a - b$ anstelle von $a + (-b)$. Das multiplikative Inverse zu $a \neq (0,0)$ notieren wir wie im Beweis als a^{-1} .

Obwohl es sich bei den Quaternionen um keinen Körper handelt, können wir aufgrund der Invertierbarkeit der Multiplikation zeigen, dass sie nullteilerfrei sind.

Proposition 6.2.4. *Die Quaternionen sind nullteilerfrei.*

Beweis. Seien $a, b \in \mathbb{H}$ mit $a \cdot b = (0,0)$. Falls $a = (0,0)$, sind wir fertig. Andernfalls gibt es ein $a^{-1} \in \mathbb{H}$ mit $a^{-1} \cdot a = (1,0)$. Multiplizieren wir die erste Gleichung von links mit a^{-1} , so erhalten wir $a^{-1} \cdot a \cdot b = (0,0)$, woraus $b = (0,0)$ folgt. $\square$

Die bekannten Kürzungsregeln für die Addition und Multiplikation folgen daraus, dass die Quaternionen einen Schiefkörper bilden.

Proposition 6.2.5. *Für $a, b, c \in \mathbb{H}$ folgt aus $a + c = b + c$, dass $a = b$ ist. Ist zudem $c \neq (0,0)$ und $a \cdot c = b \cdot c$ oder $c \cdot a = c \cdot b$, so ist ebenfalls $a = b$.*

Beweis. Die Aussage folgt direkt durch Addieren von $-c$, beziehungsweise Multiplizieren mit c^{-1} auf beiden Seiten der Gleichung. $\square$

Eine Kürzungsregel für $a, b, c \in \mathbb{H}$ mit $c \neq (0,0)$ und $a \cdot c = c \cdot b$ gibt es in den Quaternionen im Allgemeinen nicht. Das lässt sich mit Hilfe von Tabelle 6.1 zeigen: Es ist $i \cdot j = k = -j \cdot i = j \cdot (-i)$, aber $i \neq -i$.

Auch die Gleichungen $a \cdot x + b = c$ und $x \cdot a + b = c$ sind in den Quaternionen stets lösbar (der Beweis lässt sich wie in Proposition 3.2.6 führen). Da die Multiplikation nicht kommutativ ist, müssen die Gleichungen aber nicht die selbe Lösung haben.

Die Quaternionen bilden daher insgesamt eine vierdimensionale assoziative reelle Divisionsalgebra mit Einselement.

6.3 Einbettung von $\mathbb{C}$ in $\mathbb{H}$

Die komplexen Zahlen lassen sich sehr einfach in die Quaternionen einbetten. Dazu gehen wir nahezu analog wie im letzten Kapitel vor.

Proposition 6.3.1. *Sei $\varphi : \mathbb{C} \rightarrow \mathbb{H}$ mit $\varphi(a) := (a,0)$. Dann ist φ eine Einbettung.*

Beweis. Es ist

$$\varphi(a + b) = (a + b, 0) = (a, 0) + (b, 0) = \varphi(a) + \varphi(b),$$

und

$$\varphi(a \cdot b) = (a \cdot b, 0) = (a \cdot b - 0 \cdot 0, a \cdot 0 + 0 \cdot \bar{b}) = (a, 0) \cdot (b, 0) = \varphi(a) \cdot \varphi(b).$$

Um die Injektivität zu zeigen, nehmen wir an, es wäre $\varphi(a) = \varphi(b)$. Dann folgt durch komponentenweisen Vergleich direkt $a = b$. Somit ist φ eine Einbettung. $\square$

Für die Notation vereinbaren wir, dass wir in der Darstellung eines Quaternions bezüglich der Basis $\{e, i, j, k\}$ (siehe Abschnitt 6.1)

$$\lambda_1 \cdot e + \lambda_2 \cdot i + \lambda_3 \cdot j + \lambda_4 \cdot k$$

das e weglassen, also direkt

$$\lambda_1 + \lambda_2 \cdot i + \lambda_3 \cdot j + \lambda_4 \cdot k$$

schreiben (wobei $\lambda_i \in \mathbb{R}$ für $1 \leq i \leq 4$). Alternativ schreiben wir ein Quaternion auch als vierdimensionalen reellen Spalten-/Zeilenvektor dieser λ_i an, oder bleiben bei der nun gewohnten Schreibweise als Paar komplexer Zahlen.

Da die Quaternionen bezüglich der Multiplikation nicht kommutativ sind, ist im Allgemeinen $b^{-1} \cdot a \neq a \cdot b^{-1}$. Würden wir die Bruchschreibweise übertragen wollen, stellt sich daher die Frage, für welchen der beiden Ausdrücke $\frac{a}{b}$ stehen soll. Damit wir nicht mit bisherigen Konventionen brechen müssen, führen wir diese Notation auf den Quaternionen nicht allgemein ein. Es gilt allerdings folgendes Resultat.

Proposition 6.3.2. *Sei $a \in \mathbb{H}$ und $b \in \mathbb{R}$. Dann ist $a \cdot b = b \cdot a$. Es kommutieren also alle Quaternionen mit den reellen Zahlen.*

Beweis. Es entspricht b dem Quaternion $(b, 0)$, weiters ist $a = (a_1, a_2)$. Daher ist

$$a \cdot b = (a_1, a_2) \cdot (b, 0) =$$

$$\begin{aligned}
 &= (a_1 \cdot b + a_2 \cdot 0, a_1 \cdot 0 + a_2 \cdot b) = \\
 &= (b \cdot a_1 + 0 \cdot \overline{a_2}, b \cdot a_2 + 0 \cdot \overline{a_1}) = \\
 &= (b, 0) \cdot (a_1, a_2) = \\
 &= b \cdot a.
 \end{aligned}$$

□

Für $b \in \mathbb{R} \setminus \{0\}$ ist damit $b^{-1} \cdot a = a \cdot b^{-1}$. Wir vereinbaren daher für *reelle* $b \neq 0$ und $a \in \mathbb{H}$ die Schreibweise $\frac{a}{b} = a \cdot b^{-1}$.

In den komplexen Zahlen gelten die Eigenschaften $\overline{a + b} = \overline{a} + \overline{b}$, $\overline{a \cdot b} = \overline{a} \cdot \overline{b}$ und $a \cdot \overline{a} = |a|^2$. Ein ähnliches Resultat erhalten wir auch in den Quaternionen.

Proposition 6.3.3. *Seien $a, b \in \mathbb{H}$. Dann ist $\overline{a + b} = \overline{a} + \overline{b}$ und $\overline{a \cdot b} = \overline{b} \cdot \overline{a}$. Weiters ist $a \cdot \overline{a} = \overline{a} \cdot a = |a|^2$.*

Beweis. Für $a = (a_1, a_2)$ und $b = (b_1, b_2)$ ist

$$\begin{aligned}
 \overline{a + b} &= \overline{(a_1, a_2) + (b_1, b_2)} = \\
 &= \overline{(a_1 + b_1, a_2 + b_2)} = \\
 &= (\overline{a_1 + b_1}, -(a_2 + b_2)) = \\
 &= (\overline{a_1} + \overline{b_1}, -a_2 - b_2) = \\
 &= (\overline{a_1}, -a_2) + (\overline{b_1}, -b_2) = \\
 &= \overline{a} + \overline{b}.
 \end{aligned}$$

Weiters gilt für die Multiplikation

$$\begin{aligned}
 \overline{a \cdot b} &= \overline{(a_1, a_2) \cdot (b_1, b_2)} = \\
 &= \overline{(a_1 \cdot b_1 - a_2 \cdot \overline{b_2}, a_1 \cdot b_2 + a_2 \cdot \overline{b_1})} = \\
 &= (\overline{a_1 \cdot b_1 - a_2 \cdot \overline{b_2}}, -(a_1 \cdot b_2 + a_2 \cdot \overline{b_1})) = \\
 &= (\overline{a_1 \cdot b_1} - \overline{a_2 \cdot \overline{b_2}}, -a_1 \cdot b_2 - a_2 \cdot \overline{b_1}) = \\
 &= (\overline{b_1} \cdot \overline{a_1} - b_2 \cdot \overline{a_2}, -\overline{b_1} \cdot a_2 - b_2 \cdot a_1) = \\
 &= (\overline{b_1}, -b_2) \cdot (\overline{a_1}, -a_2) =
 \end{aligned}$$

$$\begin{aligned}
 &= \overline{(b_1, b_2)} \cdot \overline{(a_1, a_2)} = \\
 &= \bar{b} \cdot \bar{a}.
 \end{aligned}$$

Betrachten wir das Produkt $a \cdot \bar{a}$, so erhalten wir

$$\begin{aligned}
 a \cdot \bar{a} &= (a_1, a_2) \cdot (\bar{a}_1, -a_2) = \\
 &= (a_1 \cdot \bar{a}_1 - a_2 \cdot (-\bar{a}_2), a_1 \cdot (-a_2) + a_2 \cdot \bar{a}_1) = \\
 &= (|a_1|^2 + |a_2|^2, 0),
 \end{aligned}$$

also die reelle Zahl $|a_1|^2 + |a_2|^2 = |a|^2$. □

Betrachten wir nochmals Tabelle 6.1. Anhand dieser lässt sich leicht ablesen, dass $i^2 = j^2 = k^2 = -1$ ist. Die Gleichung $x^2 + 1 = 0$ hat in den Quaternionen demnach jedenfalls *drei* verschiedene Lösungen – obwohl sie nur vom Grad Zwei ist. Mehr noch: Wählen wir $(a, b) \in \mathbb{H}$ mit $a \in \{r \cdot i \mid r \in \mathbb{R}\}$, sodass $|a|^2 + |b|^2 = 1$ ist, so gilt wegen $a^2 = -|a|^2$ und $\bar{a} = -a$

$$(a, b)^2 = (a \cdot a - b \cdot \bar{b}, a \cdot b + b \cdot \bar{a}) = (-|a|^2 - |b|^2, 0) = (-1, 0).$$

Für die Gleichung $x^2 + 1 = 0$ gibt es in den Quaternionen somit *beliebig viele* Lösungen [20, S. 164].

6.4 Divisionsalgebren im $\mathbb{R}^{2 \cdot n+1}$, Satz von Frobenius

Wie in der Einleitung dieses Kapitels bereits beschrieben, verbrachte Hamilton viele Jahre vergeblich damit, eine Divisionsalgebra auf den dreidimensionalen reellen Zahlen $\mathbb{R}^3$ einzuführen. Der Grund für die jahrelange, letztendlich vergebliche Suche liegt darin, dass es so eine Algebra nicht gibt: Jede reelle Divisionsalgebra mit ungerader Dimension muss isomorph zu den reellen Zahlen sein.

Proposition 6.4.1 ([20, S. 155]). *Sei A eine reelle Divisionsalgebra mit ungerader Dimension und Einselement e . Dann ist A isomorph zu den reellen Zahlen.*

Beweis. Für ein $\mathbf{a} \in A$ betrachten wir die lineare Abbildung

$$f_{\mathbf{a}} : A \rightarrow A \text{ mit } f_{\mathbf{a}}(\mathbf{x}) := \mathbf{a} \cdot \mathbf{x}.$$

Da die Dimension von A ungerade ist, besitzt das charakteristische Polynom von $f_{\mathbf{a}}$ eine reelle Nullstelle $\lambda \in \mathbb{R}$. Es ist also λ ein Eigenwert von $f_{\mathbf{a}}$. Sei $\mathbf{v} \neq \mathbf{0}$ ein Eigenvektor zu λ . Dann ist

$$\mathbf{a} \cdot \mathbf{v} = \lambda \cdot \mathbf{v}, \text{ und damit } (\mathbf{a} - \lambda \cdot \mathbf{e}) \cdot \mathbf{v} = \mathbf{0}.$$

Da jede Divisionsalgebra nullteilerfrei ist, folgt daraus $\mathbf{a} = \lambda \cdot \mathbf{e}$. Es lassen sich also alle $\mathbf{a} \in A$ als Produkt einer reellen Zahl mit dem Einselement darstellen. Somit ist $A \subseteq \mathbb{R} \cdot \mathbf{e} \subseteq A$, und daher $A = \mathbb{R} \cdot \mathbf{e}$.

Wir betrachten nun die Abbildung

$$\varphi : \mathbb{R} \rightarrow \mathbb{R} \cdot \mathbf{e} \text{ mit } \varphi(\alpha) := \alpha \cdot \mathbf{e}.$$

Da $\mathbb{R} \cdot \mathbf{e}$ eindimensional ist, stellt $\{\mathbf{e}\}$ eine Basis dieses Raums dar. Aus $\varphi(\alpha) = \varphi(\beta)$, also $\alpha \cdot \mathbf{e} = \beta \cdot \mathbf{e}$, folgt wegen der Eindeutigkeit der Darstellung eines Vektors bezüglich seiner Basis $\alpha = \beta$. Die Funktion φ ist damit injektiv.

Für $\mathbf{w} \in \mathbb{R} \cdot \mathbf{e}$ gibt es ein $\alpha \in \mathbb{R}$ mit $\mathbf{w} = \alpha \cdot \mathbf{e} = \varphi(\alpha)$. Somit ist φ surjektiv, und insgesamt bijektiv. Es ist also $\mathbb{R}$ isomorph zu $\mathbb{R} \cdot \mathbf{e}$, und damit $\mathbb{R}$ isomorph zu A . $\square$

Noch ungeklärt ist die Frage, wie es mit assoziativen, endlichdimensionalen Divisionsalgebren auf $\mathbb{R}^n$ aussieht, wenn n gerade ist. Der Satz von Frobenius charakterisiert diese Algebren scharf. Im Folgenden skizzieren wir eine Beweisidee für diesen Satz nach [46; 21, S. 65–68].

Satz 6.4.2. *Jede assoziative, endlichdimensionale, reelle Divisionsalgebra D mit Einselement $\mathbf{e}$ ist isomorph zu den reellen Zahlen, zu den komplexen Zahlen, oder zu den Quaternionen.*

Beweisidee. Voraussetzung für den Beweis ist, dass jede endlichdimensionale reelle, kommutative Divisionsalgebra höchstens zweidimensional ist [20, S. 194].

Die Grundidee des Beweises beruht auf der Beobachtung, dass wir die Quaternionen als direkte Summe zweier Untervektorräume anschreiben können, die von den Einheiten erzeugt werden.

$$\mathbb{H} = \text{span}_{\mathbb{R}}\{1, i\} \oplus \text{span}_{\mathbb{R}}\{j, k\}$$

Bei $\text{span}_{\mathbb{R}}\{1, i\}$ handelt es sich um die komplexen Zahlen, und es lässt sich beweisen, dass

$$\text{span}_{\mathbb{R}}\{1, i\} = \{x \in \mathbb{H} \mid x \cdot i = i \cdot x\}$$

ist. Ebenso kann man zeigen, dass

$$\text{span}_{\mathbb{R}}\{j, k\} = \{x \in \mathbb{H} \mid x \cdot i = -i \cdot x\}$$

gilt. Da D eine reelle Divisionsalgebra ist, dürfen wir uns D als Erweiterung der reellen Zahlen vorstellen: $\mathbb{R} \subseteq D$.

Ist $D \neq \mathbb{R}$, so muss es ein $\mathbf{d} \in D \setminus \mathbb{R}$ geben. Es lässt sich beweisen, dass

$$\text{span}_{\mathbb{R}}\{\mathbf{e}, \mathbf{d}\} = \{x \in D \mid x \cdot \mathbf{d} = \mathbf{d} \cdot x\} \cong \mathbb{C} \quad (6.1)$$

erfüllt ist. Demzufolge gibt es ein $\mathbf{i} \in D$ mit $\mathbf{i}^2 = -\mathbf{e}$, und wegen Gleichung 6.1 ist $\text{span}_{\mathbb{R}}\{\mathbf{e}, \mathbf{i}\} \cong \mathbb{C}$.

Die weiteren Überlegungen basieren nun darauf, dass D wegen dieser Isomorphie als $\mathbb{C}$ -Vektorraum gesehen werden kann. Die lineare Abbildung $T : D \rightarrow D$ mit $T(\mathbf{x}) := \mathbf{x} \cdot \mathbf{i}$ besitzt die komplexen Eigenwerte $\pm i$, welche die Eigenräume

$$\begin{aligned} D^+ &= \{\mathbf{x} \in D \mid \mathbf{x} \cdot \mathbf{i} = i \cdot \mathbf{x}\}, \\ D^- &= \{\mathbf{x} \in D \mid \mathbf{x} \cdot \mathbf{i} = -i \cdot \mathbf{x}\} \end{aligned}$$

erzeugen. Es lässt sich zeigen, dass $D = D^+ \oplus D^-$ ist. Aus Gleichung 6.1 erhalten wir $D^+ \cong \mathbb{C}$. Ist nun $D^- = \emptyset$, so gilt wegen der vorigen Gleichheit $D = D^+ \cong \mathbb{C}$.

Wir nehmen daher $D^- \neq \emptyset$ an. Es lässt sich zeigen, dass $\dim_{\mathbb{R}} D^- = \dim_{\mathbb{R}} D^+ = 2$ ist, woraus wegen der Dimensionsformel, angewendet auf $D = D^+ \oplus D^-$, folgt, dass $\dim_{\mathbb{R}} D = 4$ ist.

Sei nun $\mathbf{v} \in D^-$ mit $\mathbf{v} \neq \mathbf{0}$. Dann ist $\mathbf{v} \cdot \mathbf{v} \cdot \mathbf{i} = \mathbf{v} \cdot (-\mathbf{i} \cdot \mathbf{v}) = \mathbf{i} \cdot \mathbf{v} \cdot \mathbf{v}$, und damit $\mathbf{v}^2 \in D^+$. Da $\mathbf{v} \notin \mathbb{R}$, muss nach Gleichung 6.1 $\text{span}_{\mathbb{R}}\{\mathbf{e}, \mathbf{v}\} \cong \mathbb{C}$ sein, womit $\mathbf{v}^2 \in \text{span}_{\mathbb{R}}\{\mathbf{e}, \mathbf{v}\}$ erfüllt ist.

Es ist somit $\mathbf{v}^2 \in D^+ \cap \text{span}_{\mathbb{R}}\{\mathbf{e}, \mathbf{v}\} = \text{span}_{\mathbb{R}}\{\mathbf{e}\} \cong \mathbb{R}$. Daher gibt es ein $r \in \mathbb{R}$ mit $\mathbf{v}^2 = r \cdot \mathbf{e}$. Wäre $r > 0$, so wären $\sqrt{r} \cdot \mathbf{e}$, $-\sqrt{r} \cdot \mathbf{e}$ und $\mathbf{v}$ drei verschiedene Quadratwurzeln von $\mathbf{v}^2$ – da $\text{span}_{\mathbb{R}}\{\mathbf{e}, \mathbf{v}\} \cong \mathbb{C}$ ist, kann es aber höchstens zwei verschiedene Quadratwurzeln geben, ein Widerspruch. Daher bleibt wegen $\mathbf{v} \neq \mathbf{0}$ nur, dass $r < 0$ ist.

Wir können also $\mathbf{v}^2 = r \cdot \mathbf{e}$ durch Multiplikation eines positiven reellen Faktors $s^2 \in \mathbb{R}$ so skalieren, dass $(s \cdot \mathbf{v})^2 = -\mathbf{e}$ ist. Setzen wir $\mathbf{j} := s \cdot \mathbf{v}$ und $\mathbf{k} := \mathbf{i} \cdot \mathbf{j}$, so sind $\mathbf{j}$ und $\mathbf{k}$ linear unabhängig, und daher eine Basis von D^- . Damit gilt:

$$D = \text{span}_{\mathbb{R}}\{\mathbf{e}, \mathbf{i}\} \oplus \text{span}_{\mathbb{R}}\{\mathbf{j}, \mathbf{k}\}.$$

Durch Ausrechnen der Produkte der Basisvektoren lässt sich verifizieren, dass diese mit Tabelle 6.1 übereinstimmen. Es ist somit $D \cong \mathbb{H}$. □

Bis auf Isomorphie sind daher $\mathbb{R}$, $\mathbb{C}$ und $\mathbb{H}$ die einzigen endlichdimensionalen, assoziativen, reellen Divisionsalgebren. Wollen wir eine weitere Erweiterung durchführen, müssen wir beispielsweise die Assoziativität aufgeben. Das führt im nächsten Kapitel auf die Oktonionen.

7 Die Oktonionen

Die Oktonionen wurden im Dezember 1843, also im Jahr der Entdeckung der Quaternionen, durch John Graves entdeckt. Graves hat bereits im Oktober 1843 auf einen Brief Hamiltons hin überlegt, ob sich die Konstruktionsschritte der Quaternionen nicht wiederholen ließen, und konnte Hamilton nach zwei Monaten über seine Entdeckung einer achtdimensionalen Erweiterung der reellen Zahlen schreiben, die er *Oktaven* nannte.

Unabhängig von Graves erforschte Arthur Cayley die Quaternionen und veröffentlichte im Jahr 1845 einen Artikel, in dem er die von Graves entdeckten (allerdings noch nicht publizierten) Oktonionen beschrieb. Seither werden die Oktonionen, trotz Bemühungen von Graves und Hamilton, die Priorität Graves' hervorzuheben, auch *Cayley-Zahlen* genannt [2, S. 146–147].

Während Hamilton bereits nach der Entdeckung die Anwendungen der Quaternionen für die räumliche Geometrie klar waren, war die inner- und außermathematische Bedeutung der Oktonionen lange Zeit unbekannt. Im Jahr 1983 wurde bekannt, dass sich die Oktonionen zur Beschreibung einiger Eigenschaften der Stringtheorie eignen. Innerhalb der Mathematik zeigen die Oktonionen unter anderem Resultate aus der abstrakten Algebra und Topologie auf [2, S. 147–148].

7.1 Konstruktion der Oktonionen

Die komplexen Zahlen haben wir durch Verdoppelung der reellen Zahlen als $\mathbb{R} \times \mathbb{R}$ definiert, während wir die Quaternionen durch Verdoppelung der komplexen Zahlen als $\mathbb{C} \times \mathbb{C}$ erhalten haben. Es liegt daher nahe, diese Verdoppelung erneut durchzuführen, um zu den Oktonionen zu gelangen.

Definition 7.1.1 ([2, S. 154]). Die Menge der Oktonionen $\mathbb{O}$ (auch *Oktaven*) ist definiert als

$$\mathbb{O} := \mathbb{H} \times \mathbb{H} = \{(a, b) \mid a, b \in \mathbb{H}\}.$$

Betrachten wir die Oktonionen als $\mathbb{R}$ -Vektorraum über den Quaternionen, so finden wir die Basis

$$\begin{aligned} e &:= (1, 0) & i &:= (i, 0) & j &:= (j, 0) & k &:= (k, 0) \\ l &:= (0, 1) & m &:= (0, i) & n &:= (0, j) & o &:= (0, k). \end{aligned}$$

Die Oktonionen bilden daher einen achtdimensionalen Vektorraum über den reellen Zahlen.

7.2 Rechenoperationen auf $\mathbb{O}$

Wie am Ende des letzten Kapitels bereits bemerkt, ist die Multiplikation in Erweiterungen, welche über die Quaternionen hinausgehen, nicht mehr assoziativ. Wir werden allerdings sehen, dass in den Oktonionen eine schwächere Form des Assoziativgesetzes gilt.

Definition 7.2.1. Die Addition $+$: $\mathbb{O} \times \mathbb{O} \rightarrow \mathbb{O}$ ist definiert durch

$$(a, b) + (c, d) := (a + c, b + d).$$

Die Multiplikation definieren wir wie in den Quaternionen, wobei nun wegen der fehlenden Kommutativität die Reihenfolge der Elemente relevant ist.

Definition 7.2.2 ([20, S. 211]). Die Multiplikation $\cdot$: $\mathbb{O} \times \mathbb{O} \rightarrow \mathbb{O}$ ist definiert durch

$$(a, b) \cdot (c, d) := (a \cdot c - \bar{d} \cdot b, b \cdot \bar{c} + d \cdot a).$$

Die Produkte der Basiselemente der Oktonionen sind in Tabelle 7.1 übersichtlich dargestellt.

$\cdot$	e	i	j	k	l	m	n	o
e	e	i	j	k	l	m	n	o
i	i	$-e$	k	$-j$	m	$-l$	$-o$	n
j	j	$-k$	$-e$	i	n	o	$-l$	$-m$
k	k	j	$-i$	$-e$	o	$-n$	m	$-l$
l	l	$-m$	$-n$	$-o$	$-e$	i	j	k
m	m	l	$-o$	n	$-i$	$-e$	$-k$	j
n	n	o	l	$-m$	$-j$	k	$-e$	$-i$
o	o	$-n$	m	l	$-k$	$-j$	i	$-e$

Tabelle 7.1: Multiplikationstabelle der Oktonionen [20, S. 215; 50, S. 355]

Wie bereits eingangs erwähnt, kann die Multiplikation in den Oktonionen nach dem Satz von Frobenius nicht assoziativ sein. Das folgt auch aus Tabelle 7.1:

$$(i \cdot l) \cdot o = m \cdot o = j$$

$$i \cdot (l \cdot o) = i \cdot k = -j.$$

Dennoch gilt eine speziellere Form des Assoziativgesetzes, das *Alternativgesetz*. Die Struktur der Oktonionen nennt man daher entsprechend einen Alternativkörper [50, S. 356]. Um das zu beweisen, benötigen wir folgendes Lemma.

Lemma 7.2.3. Seien $a, b \in \mathbb{H}$. Dann ist $a \cdot (b + \bar{b}) = (b + \bar{b}) \cdot a$.

Beweis. Seien $a = (a_1, a_2)$ und $b = (b_1, b_2)$ Quaternionen. Dann ist

$$b + \bar{b} = \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} + \begin{pmatrix} \bar{b}_1 \\ -b_2 \end{pmatrix} = \begin{pmatrix} b_1 + \bar{b}_1 \\ 0 \end{pmatrix}.$$

Wegen $\overline{b_1 + \bar{b}_1} = b_1 + \bar{b}_1$ und der Kommutativität in den komplexen Zahlen erhalten wir

$$\begin{aligned} a \cdot (b + \bar{b}) &= \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot \begin{pmatrix} b_1 + \bar{b}_1 \\ 0 \end{pmatrix} = \\ &= \begin{pmatrix} a_1 \cdot (b_1 + \bar{b}_1) \\ a_2 \cdot \overline{(b_1 + \bar{b}_1)} \end{pmatrix} = \end{aligned}$$

$$\begin{aligned}
&= \begin{pmatrix} a_1 \cdot (b_1 + \overline{b_1}) \\ a_2 \cdot (b_1 + \overline{b_1}) \end{pmatrix} = \\
&= \begin{pmatrix} (b_1 + \overline{b_1}) \cdot a_1 \\ (b_1 + \overline{b_1}) \cdot a_2 \end{pmatrix} = \\
&= \begin{pmatrix} b_1 + \overline{b_1} \\ 0 \end{pmatrix} \cdot \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} = \\
&= (b + \overline{b}) \cdot a.
\end{aligned}$$

□

Proposition 7.2.4. $(\mathbb{O}, +, \cdot, (0,0), (1,0))$ ist ein Alternativkörper. Es sind daher folgende Eigenschaften erfüllt:

- (1) $\forall a, b, c \in \mathbb{O} : (a + b) + c = a + (b + c)$ (Assoziativität von $+$).
- (2) $\forall a, b \in \mathbb{O} : a + b = b + a$ (Kommutativität von $+$).
- (3) $\forall a \in \mathbb{O} : a + (0,0) = (0,0)$ (Nullelement).
- (4) $\forall a \in \mathbb{O} : \exists -a \in \mathbb{O} : a + (-a) = (0,0)$ (Invertierbarkeit von $+$).
- (5) $\forall a, b \in \mathbb{O} : (a \cdot a) \cdot b = a \cdot (a \cdot b)$ und $(a \cdot b) \cdot b = a \cdot (b \cdot b)$ (Alternativität von $\cdot$).
- (6) $\forall a \in \mathbb{O} : a \cdot (1,0) = (1,0) \cdot a = a$ (Einselement).
- (7) $\forall a \in \mathbb{O} \setminus \{(0,0)\} : \exists a^{-1} \in \mathbb{O} : a \cdot a^{-1} = a^{-1} \cdot a = (1,0)$ (Invertierbarkeit von $\cdot$).
- (8) $\forall a, b, c \in \mathbb{O} : a \cdot (b + c) = a \cdot b + a \cdot c$ und $(b + c) \cdot a = b \cdot a + c \cdot a$ (Distributivität).

Beweis. Seien $a = (a_1, a_2)$, $b = (b_1, b_2)$ und $c = (c_1, c_2)$ Oktonionen. Die Eigenschaften (1) bis (4) folgen direkt daraus, dass die Addition komponentenweise definiert ist, und die entsprechenden Eigenschaften in den Quaternionen erfüllt sind.

Zur Alternativität der Multiplikation: Zwecks besserer Lesbarkeit schreiben wir Oktonionen als Spaltenvektoren. Es ist dann

$$\begin{aligned}
(a \cdot a) \cdot b &= \left(\begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \right) \cdot \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} = \\
&= \begin{pmatrix} a_1 \cdot a_1 - \overline{a_2} \cdot a_2 \\ a_2 \cdot \overline{a_1} + a_2 \cdot a_1 \end{pmatrix} \cdot \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} =
\end{aligned}$$

$$\begin{aligned}
 &= \left((a_1 \cdot a_1 - \overline{a_2} \cdot a_2) \cdot b_1 - \overline{b_2} \cdot (a_2 \cdot \overline{a_1} + a_2 \cdot a_1) \right) = \\
 &= \left((a_2 \cdot \overline{a_1} + a_2 \cdot a_1) \cdot \overline{b_1} + b_2 \cdot (a_1 \cdot a_1 - \overline{a_2} \cdot a_2) \right) = \\
 &= \left(a_1 \cdot a_1 \cdot b_1 - \overline{a_2} \cdot a_2 \cdot b_1 - \overline{b_2} \cdot a_2 \cdot \overline{a_1} - \overline{b_2} \cdot a_2 \cdot a_1 \right) = \\
 &= \left(a_2 \cdot \overline{a_1} \cdot \overline{b_1} + a_2 \cdot a_1 \cdot \overline{b_1} + b_2 \cdot a_1 \cdot a_1 - b_2 \cdot \overline{a_2} \cdot a_2 \right).
 \end{aligned}$$

Aus Proposition 6.3.3 folgt, dass $\overline{a_2} \cdot a_2$ eine reelle Zahl ist, die nach Proposition 6.3.2 mit den Quaternionen kommutiert. Es ist daher

$$\overline{a_2} \cdot a_2 \cdot b_1 = b_1 \cdot \overline{a_2} \cdot a_2 \text{ und } b_2 \cdot \overline{a_2} \cdot a_2 = a_2 \cdot \overline{a_2} \cdot b_2.$$

Wegen Lemma 7.2.3 ist

$$\begin{aligned}
 \overline{b_2} \cdot a_2 \cdot \overline{a_1} + \overline{b_2} \cdot a_2 \cdot a_1 &= \overline{a_1} \cdot \overline{b_2} \cdot a_2 + a_1 \cdot \overline{b_2} \cdot a_2, \text{ und} \\
 \overline{b_1} \cdot a_1 + \overline{b_1} \cdot \overline{a_1} &= a_1 \cdot \overline{b_1} + \overline{a_1} \cdot \overline{b_1}.
 \end{aligned}$$

Multiplizieren wir die zweite Gleichung von links mit a_2 , so erhalten wir

$$a_2 \cdot \overline{b_1} \cdot a_1 + a_2 \cdot \overline{b_1} \cdot \overline{a_1} = a_2 \cdot a_1 \cdot \overline{b_1} + a_2 \cdot \overline{a_1} \cdot \overline{b_1}.$$

Wir können nun die obige Gleichungskette weiterführen:

$$\begin{aligned}
 (a \cdot a) \cdot b &= \left(a_1 \cdot a_1 \cdot b_1 - \overline{a_2} \cdot a_2 \cdot b_1 - \overline{b_2} \cdot a_2 \cdot \overline{a_1} - \overline{b_2} \cdot a_2 \cdot a_1 \right) = \\
 &= \left(a_1 \cdot a_1 \cdot b_1 - b_1 \cdot \overline{a_2} \cdot a_2 - \overline{a_1} \cdot \overline{b_2} \cdot a_2 - a_1 \cdot \overline{b_2} \cdot a_2 \right) = \\
 &= \left(a_2 \cdot \overline{b_1} \cdot a_1 + a_2 \cdot \overline{b_1} \cdot \overline{a_1} + b_2 \cdot a_1 \cdot a_1 - a_2 \cdot \overline{a_2} \cdot b_2 \right) = \\
 &= \left(a_1 \cdot (a_1 \cdot b_1 - \overline{b_2} \cdot a_2) - (b_1 \cdot \overline{a_2} + \overline{a_1} \cdot \overline{b_2}) \cdot a_2 \right) = \\
 &= \left(a_1 \cdot \overline{(a_1 \cdot b_1 - \overline{b_2} \cdot a_2)} - \overline{(a_2 \cdot \overline{b_1} + b_2 \cdot a_1)} \cdot a_2 \right) = \\
 &= \left(a_1 \cdot (a_1 \cdot b_1 - \overline{b_2} \cdot a_2) + (a_2 \cdot \overline{b_1} + b_2 \cdot a_1) \cdot a_1 \right) = \\
 &= \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot \begin{pmatrix} a_1 \cdot b_1 - \overline{b_2} \cdot a_2 \\ a_2 \cdot \overline{b_1} + b_2 \cdot a_1 \end{pmatrix} = \\
 &= a \cdot (a \cdot b).
 \end{aligned}$$

Das zweite Alternativgesetz lässt sich analog beweisen. Um zu zeigen, dass es sich bei $(1,0)$ um das Einselement der Multiplikation handelt, betrachten wir die Gleichungen

$$\begin{aligned}(a_1, a_2) \cdot (1, 0) &= (a_1 \cdot 1 - \bar{0} \cdot a_2, a_2 \cdot \bar{1} + 0 \cdot a_1) = (a_1, a_2), \\ (1, 0) \cdot (a_1, a_2) &= (1 \cdot a_1 - \bar{a_2} \cdot 0, 0 \cdot \bar{a_1} + a_2 \cdot 1) = (a_1, a_2).\end{aligned}$$

Für den Beweis der Invertierbarkeit der Multiplikation nehmen wir an, dass $a \neq (0,0)$ ist. Wie im Beweis für die Quaternionen setzen wir

$$a^{-1} = \left(\frac{\bar{a_1}}{|a_1|^2 + |a_2|^2}, \frac{-a_2}{|a_1|^2 + |a_2|^2} \right).$$

Wir erhalten dann

$$\begin{aligned}a \cdot a^{-1} &= (a_1, a_2) \cdot \left(\frac{\bar{a_1}}{|a_1|^2 + |a_2|^2}, \frac{-a_2}{|a_1|^2 + |a_2|^2} \right) = \\ &= \left(a_1 \cdot \frac{\bar{a_1}}{|a_1|^2 + |a_2|^2} - \frac{-a_2}{|a_1|^2 + |a_2|^2} \cdot a_2, a_2 \cdot \frac{\bar{a_1}}{|a_1|^2 + |a_2|^2} + \frac{-a_2}{|a_1|^2 + |a_2|^2} \cdot a_1 \right) = \\ &= \left(\frac{a_1 \cdot \bar{a_1}}{|a_1|^2 + |a_2|^2} - \frac{-\bar{a_2} \cdot a_2}{|a_1|^2 + |a_2|^2}, \frac{a_2 \cdot a_1}{|a_1|^2 + |a_2|^2} + \frac{-a_2 \cdot a_1}{|a_1|^2 + |a_2|^2} \right) = \\ &= \left(\frac{|a_1|^2}{|a_1|^2 + |a_2|^2} + \frac{|a_2|^2}{|a_1|^2 + |a_2|^2}, \frac{a_2 \cdot a_1}{|a_1|^2 + |a_2|^2} - \frac{a_2 \cdot a_1}{|a_1|^2 + |a_2|^2} \right) = \\ &= (1, 0).\end{aligned}$$

Weiters gilt

$$\begin{aligned}a^{-1} \cdot a &= \left(\frac{\bar{a_1}}{|a_1|^2 + |a_2|^2}, \frac{-a_2}{|a_1|^2 + |a_2|^2} \right) \cdot (a_1, a_2) = \\ &= \left(\frac{\bar{a_1}}{|a_1|^2 + |a_2|^2} \cdot a_1 - \bar{a_2} \cdot \frac{-a_2}{|a_1|^2 + |a_2|^2}, \frac{-a_2}{|a_1|^2 + |a_2|^2} \cdot \bar{a_1} + a_2 \cdot \frac{\bar{a_1}}{|a_1|^2 + |a_2|^2} \right) = \\ &= \left(\frac{|a_1|^2}{|a_1|^2 + |a_2|^2} + \frac{|a_2|^2}{|a_1|^2 + |a_2|^2}, \frac{-a_2 \cdot \bar{a_1}}{|a_1|^2 + |a_2|^2} + \frac{a_2 \cdot \bar{a_1}}{|a_1|^2 + |a_2|^2} \right) = \\ &= (1, 0).\end{aligned}$$

Um die Distributivgesetze zu beweisen, betrachten wir

$$\begin{aligned}
 (a_1, a_2) \cdot ((b_1, b_2) + (c_1, c_2)) &= (a_1, a_2) \cdot (b_1 + c_1, b_2 + c_2) = \\
 &= (a_1 \cdot (b_1 + c_1) - \overline{(b_2 + c_2)} \cdot a_2, a_2 \cdot (b_1 + c_1) + (b_2 + c_2) \cdot a_1) = \\
 &= (a_1 \cdot b_1 - \overline{b_2} \cdot a_2 + a_1 \cdot c_1 - \overline{c_2} \cdot a_2, a_2 \cdot \overline{b_1} + b_2 \cdot a_1 + a_2 \cdot \overline{c_1} + c_2 \cdot a_1) = \\
 &= (a_1 \cdot b_1 - \overline{b_2} \cdot a_2, a_2 \cdot \overline{b_1} + b_2 \cdot a_1) + (a_1 \cdot c_1 - \overline{c_2} \cdot a_2, a_2 \cdot \overline{c_1} + c_2 \cdot a_1) = \\
 &= (a_1, a_2) \cdot (b_1, b_2) + (a_1, a_2) \cdot (c_1, c_2),
 \end{aligned}$$

wobei wir für den Beweis die Distributivität der Quaternionen verwendet haben. Analog folgt das zweite Distributivgesetz:

$$\begin{aligned}
 ((b_1, b_2) + (c_1, c_2)) \cdot (a_1, a_2) &= (b_1 + c_1, b_2 + c_2) \cdot (a_1, a_2) = \\
 &= ((b_1 + c_1) \cdot a_1 - \overline{a_2} \cdot (b_2 + c_2), (b_2 + c_2) \cdot \overline{a_1} + a_2 \cdot (b_1 + c_1)) = \\
 &= (b_1 \cdot a_1 - \overline{a_2} \cdot b_2 + c_1 \cdot a_1 - \overline{a_2} \cdot c_2, b_2 \cdot \overline{a_1} + a_2 \cdot b_1 + c_2 \cdot \overline{a_1} + a_2 \cdot c_1) = \\
 &= (b_1 \cdot a_1 - \overline{a_2} \cdot b_2, b_2 \cdot \overline{a_1} + a_2 \cdot b_1) + (c_1 \cdot a_1 - \overline{a_2} \cdot c_2, c_2 \cdot \overline{a_1} + a_2 \cdot c_1) = \\
 &= (b_1, b_2) \cdot (a_1, a_2) + (c_1, c_2) \cdot (a_1, a_2). \quad \square
 \end{aligned}$$

Wir schreiben wieder $a - b$ statt $a + (-b)$, und a^{-1} für das multiplikative Inverse von $a \neq (0,0)$.

Aus der Alternativität der Multiplikation erhalten wir $a \cdot (a \cdot a) = (a \cdot a) \cdot a$ für alle $a \in \mathbb{O}$. Wir dürfen daher Potenzen auf den Oktonionen weiter wie gewohnt verwenden; insbesondere ist $x^m \cdot x^n = x^{m+n}$ für alle $x \in \mathbb{O}$ und $m, n \in \mathbb{N} \setminus \{0\}$. Eine Algebra, welche diese Eigenschaft erfüllt, wird *potenz-assoziativ* genannt [20, S. 151].

Obwohl die Assoziativität in den Oktonionen nicht erfüllt ist, erhalten wir folgendes Resultat.

Proposition 7.2.5 ([55, S. 171]). *Seien $a, b \in \mathbb{O}$ mit $b \neq (0,0)$. Dann ist $(a \cdot b) \cdot b^{-1} = a \cdot (b \cdot b^{-1}) = a$ und $b^{-1} \cdot (b \cdot a) = (b^{-1} \cdot b) \cdot a = a$.*

Beweis. Wir schreiben $a = (a_1, a_2)$ und $b = (b_1, b_2)$. Nach Proposition 7.2.4 ist

$$b^{-1} = \left(\frac{\overline{b_1}}{|b_1|^2 + |b_2|^2}, \frac{-b_2}{|b_1|^2 + |b_2|^2} \right).$$

Damit erhalten wir (zur besseren Lesbarkeit verwenden wir Spaltenvektoren):

$$\begin{aligned} (a \cdot b) \cdot b^{-1} &= \begin{pmatrix} a_1 \cdot b_1 - \overline{b_2} \cdot a_2 \\ a_2 \cdot \overline{b_1} + b_2 \cdot a_1 \end{pmatrix} \cdot \begin{pmatrix} \frac{\overline{b_1}}{|b_1|^2 + |b_2|^2} \\ \frac{-b_2}{|b_1|^2 + |b_2|^2} \end{pmatrix} = \\ &= \begin{pmatrix} (a_1 \cdot b_1 - \overline{b_2} \cdot a_2) \cdot \frac{\overline{b_1}}{|b_1|^2 + |b_2|^2} - \frac{-b_2}{|b_1|^2 + |b_2|^2} \cdot (a_2 \cdot \overline{b_1} + b_2 \cdot a_1) \\ (a_2 \cdot \overline{b_1} + b_2 \cdot a_1) \cdot \frac{\overline{b_1}}{|b_1|^2 + |b_2|^2} + \frac{-b_2}{|b_1|^2 + |b_2|^2} \cdot (a_1 \cdot b_1 - \overline{b_2} \cdot a_2) \end{pmatrix} = \\ &= \begin{pmatrix} (a_1 \cdot b_1 - \overline{b_2} \cdot a_2) \cdot \frac{\overline{b_1}}{|b_1|^2 + |b_2|^2} + \frac{\overline{b_2}}{|b_1|^2 + |b_2|^2} \cdot (a_2 \cdot \overline{b_1} + b_2 \cdot a_1) \\ (a_2 \cdot \overline{b_1} + b_2 \cdot a_1) \cdot \frac{\overline{b_1}}{|b_1|^2 + |b_2|^2} - \frac{b_2}{|b_1|^2 + |b_2|^2} \cdot (a_1 \cdot b_1 - \overline{b_2} \cdot a_2) \end{pmatrix} = \\ &= \frac{1}{|b_1|^2 + |b_2|^2} \cdot \begin{pmatrix} (a_1 \cdot b_1 - \overline{b_2} \cdot a_2) \cdot \overline{b_1} + \overline{b_2} \cdot (a_2 \cdot \overline{b_1} + b_2 \cdot a_1) \\ (a_2 \cdot \overline{b_1} + b_2 \cdot a_1) \cdot \overline{b_1} - b_2 \cdot (a_1 \cdot b_1 - \overline{b_2} \cdot a_2) \end{pmatrix} = \\ &= \frac{1}{|b_1|^2 + |b_2|^2} \cdot \begin{pmatrix} a_1 \cdot b_1 \cdot \overline{b_1} + \overline{b_2} \cdot b_2 \cdot a_1 \\ a_2 \cdot \overline{b_1} \cdot \overline{b_1} + b_2 \cdot \overline{b_2} \cdot a_2 \end{pmatrix} = \\ &= \frac{1}{|b_1|^2 + |b_2|^2} \cdot \begin{pmatrix} a_1 \cdot |b_1|^2 + |b_2|^2 \cdot a_1 \\ a_2 \cdot |b_1|^2 + |b_2|^2 \cdot a_2 \end{pmatrix} = \\ &= \frac{1}{|b_1|^2 + |b_2|^2} \cdot \begin{pmatrix} a_1 \cdot |b_1|^2 + a_1 \cdot |b_2|^2 \\ a_2 \cdot |b_1|^2 + a_2 \cdot |b_2|^2 \end{pmatrix} = \\ &= \begin{pmatrix} a_1 \\ a_2 \end{pmatrix}, \end{aligned}$$

wobei wir verwendet haben, dass für $a \in \mathbb{H}$ stets $a \cdot \overline{a} = |a|^2$ ist, und die Quaternionen mit den reellen Zahlen kommutieren (Proposition 6.3.2).

Die zweite Aussage lässt sich analog beweisen. □

Zwar bilden die Oktonionen keinen Körper, sind aber nullteilerfrei.

Proposition 7.2.6. *Die Oktonionen sind nullteilerfrei.*

Beweis. Seien $a, b \in \mathbb{O}$ mit $a \cdot b = (0,0)$. Falls $a = (0,0)$ ist, so sind wir fertig. Andernfalls gibt es ein a^{-1} mit $a^{-1} \cdot a = (1,0)$. Es ist dann

$$(0,0) = a^{-1} \cdot (0,0) = a^{-1} \cdot (a \cdot b) = (a^{-1} \cdot a) \cdot b = (1,0) \cdot b = b,$$

wobei wir für die Umformungen Proposition 7.2.5 verwendet haben. $\square$

Auch die bekannten Kürzungsregeln behalten wegen der Existenz von additiven und multiplikativen Inversen ihre Gültigkeit.

Proposition 7.2.7. *Für $a, b, c \in \mathbb{O}$ folgt aus $a + c = b + c$, dass $a = b$ ist. Ist zudem $c \neq (0,0)$ und $a \cdot c = b \cdot c$ oder $c \cdot a = c \cdot b$, so ist ebenfalls $a = b$.*

Beweis. Aus $a + c = b + c$ folgt durch Addieren von $-c$ auf beiden Seiten der Gleichung direkt $a = b$.

Für die zweite Aussage nehmen wir an, es wäre $c \neq (0,0)$ und $a \cdot c = b \cdot c$. Dann gibt es ein $c^{-1} \in \mathbb{O}$ mit $c \cdot c^{-1} = (1,0)$. Unter Verwendung von Proposition 7.2.5 erhalten wir

$$(a \cdot c) \cdot c^{-1} = (b \cdot c) \cdot c^{-1} \Leftrightarrow a \cdot (c \cdot c^{-1}) = b \cdot (c \cdot c^{-1}) \Leftrightarrow a = b. \quad \square$$

Die Gleichungen $a \cdot x + b = c$ und $x \cdot a + b = c$ sind in den Oktonionen für $a \neq (0,0)$ ebenfalls stets lösbar (der Beweis läuft analog wie in den rationalen Zahlen), wenn sie auch wegen der fehlenden Kommutativität nicht die selbe Lösung $x \in \mathbb{O}$ besitzen müssen. Die Oktonionen bilden somit eine alternative reelle Divisionsalgebra mit Einselement, die jedoch weder kommutativ noch assoziativ ist.

7.3 Einbettung von $\mathbb{H}$ in $\mathbb{O}$

Wir zeigen nun, dass es sich bei den Oktonionen tatsächlich um eine Erweiterung der Quaternionen handelt.

Proposition 7.3.1. *Sei $\varphi : \mathbb{H} \rightarrow \mathbb{O}$ mit $\varphi(a) := (a, 0)$. Dann ist φ eine Einbettung.*

Beweis. Für $a, b \in \mathbb{H}$ ist

$$\varphi(a + b) = (a + b, 0) = (a, 0) + (b, 0) = \varphi(a) + \varphi(b).$$

Weiters ist

$$\varphi(a \cdot b) = (a \cdot b, 0) = (a \cdot b - \bar{0} \cdot 0, 0 \cdot \bar{b} + 0 \cdot a) = (a, 0) \cdot (b, 0) = \varphi(a) \cdot \varphi(b).$$

Die Injektivität folgt direkt durch komponentenweisen Vergleich von $\varphi(a) = \varphi(b)$. $\square$

Wie in den vorigen Kapiteln vereinbaren wir, dass ein Oktonion in seiner Darstellung bezüglich der Basis $\{e, i, j, k, l, m, n, o\}$

$$\lambda_1 \cdot e + \lambda_2 \cdot i + \lambda_3 \cdot j + \lambda_4 \cdot k + \lambda_5 \cdot l + \lambda_6 \cdot m + \lambda_7 \cdot n + \lambda_8 \cdot o$$

auch als

$$\lambda_1 + \lambda_2 \cdot i + \lambda_3 \cdot j + \lambda_4 \cdot k + \lambda_5 \cdot l + \lambda_6 \cdot m + \lambda_7 \cdot n + \lambda_8 \cdot o$$

angeschrieben werden kann, wir also anstelle von $\lambda_1 \cdot e$ kürzer λ_1 schreiben dürfen.

In Satz 6.4.2 haben wir bewiesen, dass jede assoziative, endlichdimensionale, reelle Divisionsalgebra mit Einselement isomorph zu $\mathbb{R}$, $\mathbb{C}$ oder $\mathbb{H}$ ist. Ersetzen wir in dieser Aussage die Forderung nach Assoziativität durch Alternativität, erhalten wir folgenden Satz, für den wir eine Beweisidee nach [20, S. 216] skizzieren.

Satz 7.3.2 ([20, S. 216]). *Jede alternative, endlichdimensionale, reelle Divisionsalgebra D mit Einselement $\mathbf{e}$ ist isomorph zu den reellen Zahlen, den komplexen Zahlen, den Quaternionen oder zu den Oktonionen.*

Beweisidee. Angenommen D wäre assoziativ, dann würde aus Satz 6.4.2 folgen, dass D isomorph zu $\mathbb{R}$, $\mathbb{C}$ oder $\mathbb{H}$ ist. Sei daher D nichtassoziativ. Es lässt sich zeigen, dass sich dann eine echte assoziative Unteralgebra H von D mit $\mathbf{e} \in H$ finden lässt, die über eine Abbildung $\varphi : H \rightarrow \mathbb{H}$ zu den Quaternionen isomorph ist. Weiters gibt es ein $\mathbf{o} \in D$ mit $\mathbf{o}^2 = -e$, sodass $\mathbf{o}$ auf alle Vektoren in H normal steht.

Die Idee ist nun, dass jenes $\mathbf{o}$ eine imaginäre Einheit von D darstellt, ähnlich wie o eine imaginäre Einheit der Oktonionen ist, die auf alle Einheiten $\{e, i, j, k\}$ der Quaternionen normal steht. Während wir die Oktonionen als direkte Summe

$$\mathbb{O} = \mathbb{H} \oplus \mathbb{H} \cdot o$$

anschreiben können (vergleiche hierfür mit Tabelle 7.1), können wir D als Summe

$$D = H \oplus H \cdot \mathbf{o}$$

darstellen. Es lässt sich nun zeigen, dass D über die Abbildung $\psi : D \rightarrow \mathbb{O}$ mit

$$\psi(\mathbf{v} + \mathbf{w} \cdot \mathbf{o}) = \varphi(\mathbf{v}) + \varphi(\mathbf{w}) \cdot o$$

tatsächlich isomorph zu den Oktonionen ist. □

Bemerkungen zur Didaktik der Zahlenbereichserweiterungen

Der in dieser Diplomarbeit behandelte formal-axiomatische Zugang zu den Zahlenbereichen ist selbstverständlich nicht für den Schulunterricht geeignet; das beweist nicht zuletzt das Scheitern von *New Math* in der zweiten Hälfte des 20. Jahrhunderts. Dennoch sind gerade die Mengen der natürlichen, ganzen und rationalen Zahlen aus der heutigen Alltagswelt nicht mehr wegzudenken, und sind fest im Lehrplan der AHS Unterstufe verankert [44, S. 6–7]. Es stellt sich damit die Frage, wie Zahlenbereichserweiterungen im Schulunterricht prinzipiell motiviert und durchgeführt werden können.

Wie in dieser Arbeit gezeigt wurde, motiviert innermathematisch vor allem die Suche nach Lösungen von Problemen zu Erweiterungen des Zahlbegriffs; so wurden wir etwa durch die Frage nach der Lösbarkeit bestimmter Gleichungen auf die Mengen $\mathbb{Z}$, $\mathbb{Q}$ und $\mathbb{C}$ geführt, während wir über die Suche nach einer vollständigen Körpererweiterung die reellen Zahlen erschlossen haben. Auch im Schulunterricht lassen sich die Zahlenbereichserweiterungen anhand von Problemen motivieren, die sich innerhalb der Grenzen des bisherigen Bereichs nicht mehr lösen lassen. Hefendehl-Hebeker und Prediger nennen für die Erweiterungen auf $\mathbb{Z}^-$ und $\mathbb{Q}$ das Zählen und Messen als Beispiel: Durch die Erweiterung auf die rationalen Zahlen lässt sich die Genauigkeit beim Teilen erhöhen, während es die Erweiterung auf die negativen Zahlen ermöglicht, Größen relativ zu einer festen Marke zu bestimmen (Temperatur, Saldo) [27, S. 2].

Die Erweiterung auf die (irrationalen) reellen Zahlen lässt sich eher weniger sinnvoll anhand praktischer Probleme motivieren. Weder spielen irrationale Zahlen im Alltagsleben eine mit den anderen Zahlenbereichen vergleichbare herausragende Rolle, noch erscheinen Eigenschaften der reellen Zahlen als besonders. In Abschnitt 4.2 haben wir

die Notwendigkeit einer Erweiterung damit argumentiert, dass bestimmte Beispiele, die anschaulich eine Lösung besitzen, innerhalb der rationalen Zahlen nicht lösbar sind. Innermathematisch hat uns das auf den Begriff der Vollständigkeit gebracht, der sich am ehesten als „Lückenlosigkeit des Zahlenstrahls“ übersetzen lässt. Die Vollständigkeit wird im Unterricht aber in der Regel nicht explizit herausgearbeitet [25, S. 36] – nicht zuletzt deswegen, weil es sich bei ihr um eine selbstverständliche Eigenschaft handelt, deren Existenz kaum Schülerinnen und Schüler bezweifeln würden. Wohl aber lässt sich sehr elegant beweisen, dass die rationalen Zahlen nicht ausreichen, um sämtliche (geometrischen) Größen zu beschreiben, und wir daher unser *Verständnis* des Zahlbegriffs erweitern müssen (Proposition 4.2.1 kann in der Form auch der Schule bewiesen werden). Möglichkeiten, sich im Unterricht mit der Vollständigkeit und Irrationalität zu beschäftigen, finden sich etwa in [16].

Die Erweiterung auf die komplexen Zahlen haben wir innermathematisch mit dem Problem motiviert, Nullstellen von Polynomen wie $x^2 + 1$ zu finden. Dieser Weg lässt sich auch im Unterricht nachgehen, wo in der Regel $i := \sqrt{-1}$ definiert wird. Bei der Einführung der imaginären Einheit ist jedenfalls darauf zu achten, dass den Schülerinnen und Schülern nicht der Eindruck von Willkürlichkeit in der Mathematik vermittelt wird: Neue Einheiten lassen sich nicht nach Lust und Laune einführen, der neue Zahlenbereich soll möglichst viele Eigenschaften des vorigen Zahlenbereichs übernehmen (Permanenzprinzip). Thies gibt ein kurzes, aber einfaches Beispiel dafür an, wohin reine Willkür bei Zahlenbereichserweiterungen führen kann [52, S. 58]:

Sei j die Lösung der in $\mathbb{R}$ nicht lösbaren Gleichung $0 \cdot x = 1$. Man erhält unter Berücksichtigung der Rechengesetze

$$0 \cdot j = (0 + 0) \cdot j = 0 \cdot j + 0 \cdot j,$$

also $0 \cdot j = 0$. Andererseits ist nach Voraussetzung $0 \cdot j = 1$, also folgt $0 = 1$, ein Widerspruch.

Im Mathematik-Lehrplan der AHS stellen die komplexen Zahlen leider ein relativ isoliertes Thema dar, das seinen Höhepunkt in der Regel im Fundamentalsatz der Algebra findet. Um die Anwendungsmöglichkeiten der komplexen Zahlen zu zeigen, ist ein fachübergreifender Unterricht mit Physik sinnvoll. Dort lassen sich komplexe Zahlen etwa

in der Elektrotechnik anwenden. Innermathematisch könnte in der 12. Schulstufe, zum Abschluss des Kapitels über dynamische Prozesse, eine Einheit zum Newton-Verfahren auf $\mathbb{C}$ und Fraktalen durchgeführt werden [52, S. 61].

Erweiterungen auf die Quaternionen und Oktonionen sind im Lehrplan der AHS nicht vorgesehen [44; 43]. Sobald im Unterricht vermittelt wurde, dass sich komplexe Zahlen auch als zweidimensionale Koordinaten in der Gaußschen Zahlenebene ansehen lassen, bieten sich allerdings Überlegungen an, ob man nicht auch auf dem aus der analytischen Geometrie bekannten dreidimensionalen Raum ebenfalls Addition und Multiplikation einführen könnte. Mit Hilfe eines einfachen Arguments, welches ohne Kenntnisse über lineare Algebra auskommt, lässt sich zeigen, dass sich auf dem $\mathbb{R}^3$ jedenfalls keine Multiplikation mit bekannten Eigenschaften definieren lässt [38].

Als Abschluss des Kapitels zu den komplexen Zahlen kann im Unterricht ein Überblick über die bisherigen Zahlenbereiche gegeben werden. Folgende Punkte könnten dabei thematisiert werden:

- Darstellung der Zahlenmengen als Mengendiagramm ($\mathbb{N} \subseteq \mathbb{Z} \subseteq \dots$).
- Welche Gründe gab es für die jeweiligen Erweiterungen?
- Welche Eigenschaften wurden bei einer Erweiterung gewonnen, welche verloren?

An dieser Stelle lässt sich thematisieren, dass auch die komplexen Zahlen zu einem weiteren Zahlenbereich erweitert werden können, und dieser wiederum ebenfalls. Es sollte allerdings hervorgehoben werden, dass mit jeder Erweiterung nützliche Eigenschaften der Zahlenbereiche verloren gehen, und sie somit immer schwieriger zu handhaben sind.

Bei der Durchführung der Zahlenbereichserweiterungen im Unterricht sollte jedenfalls bedacht werden, dass diese den Schülerinnen und Schülern oftmals schwer fällt. Als Grund dafür zitieren Hefendehl-Hebeker und Prediger mehrere Forschungsprojekte, denen zufolge diese Probleme an Wandlungen in den Vorstellungen zu den Zahlenbereichen liegen [27, S. 2]. So gilt etwa die aus der Volksschule vertraute Eigenschaft, dass das Produkt zweier Zahlen stets größer gleich den einzelnen Faktoren ist, bereits beim Übergang zu den Dezimalzahlen nicht mehr. Hinzu kommt bei der Erweiterung auf die negativen Zahlen die Zweideutigkeit des Minuszeichens, das nun sowohl für

einen Rechenoperator, als auch für ein Vorzeichen steht. Die komplexen Zahlen brechen schließlich mit der vertrauten Schreibweise: hier sind nun *zwei* Zahlen zur Beschreibung *einer* Zahl notwendig.

Zur Lösung dieser Probleme schlagen Hefendehl-Hebeker und Prediger unter anderem vor, diese Stolpersteine explizit zu thematisieren, und sie nicht didaktisch zu umschiffen [27, S. 5–7]. Malle macht den Vorschlag, Schwierigkeiten in der historischen Entstehung der Zahlenbereiche auf ihre Relevanz für den Schulunterricht zu untersuchen, da manche der historischen Probleme auch von Schülerinnen und Schülern nachgegangen werden (etwa das Anerkennen neuer Zahlenbereiche) [37].

Ausblick

Ausgehend von den reellen Zahlen haben wir, jeweils durch Verdoppelung des vorigen Zahlenbereiches, die komplexen Zahlen, die Quaternionen und Oktonionen als zwei-, vier-, beziehungsweise achtdimensionale Divisionsalgebren über den reellen Zahlen erhalten. Es liegt daher nahe, dieses Verdoppelungsverfahren erneut durchzuführen, um einen weiteren Zahlenbereich, eine sechzehndimensionale reelle Algebra, zu erschließen.

Betrachten wir allerdings die obigen Zahlenbereiche genauer, so können wir feststellen, dass wir für jede der Erweiterungen einen Preis zahlen mussten: Die komplexen Zahlen lassen sich nicht anordnen, in den Quaternionen ist die Multiplikation nicht mehr kommutativ, und in den Oktonionen zusätzlich auch nicht assoziativ. Wir würden daher vermuten, dass wir auch für die Erweiterung auf jene sechzehndimensionale Algebra, die *Sedenionen*, eine besondere Eigenschaft aufgeben müssten – und tatsächlich ist das so: Die Sedenionen sind nicht mehr alternativ und besitzen Nullteiler, weswegen sie keine Divisionsalgebra bilden. Damit ist bereits das Lösen von einfachen linearen Gleichungen aufwändig [32, S. 79–80, 82–84].

Das Verdoppelungsverfahren liefert uns 2^n -dimensionale reelle Algebren, von denen wir aufgrund obiger Überlegungen vermuten würden, dass sie beginnend mit den Sedenionen eher uninteressant sind. Wir können uns nun ähnlich wie Hamilton die Frage stellen, ob es nicht interessante Divisionsalgebren in anderen Dimensionen gibt. Wegen Proposition 6.4.1 müsste die Dimension einer solchen Algebra gerade sein; unterhalb von Sechzehn blieben also genau $\mathbb{R}^6$, $\mathbb{R}^{10}$, $\mathbb{R}^{12}$ und $\mathbb{R}^{14}$ als mögliche Kandidaten übrig. Heinz Hopf konnte allerdings zeigen, dass diese Fälle nicht möglich sind: Die Dimension einer reellen Divisionsalgebra entspricht stets einer Zweierpotenz [20, S. 234].

Es bleibt somit die Frage, ob wir uns mit unserer Vermutung geirrt haben, dass den Zahlenbereichen ab den Oktonionen zunehmend wichtige Eigenschaften verloren gehen. Könnte nicht beispielsweise auf dem $\mathbb{R}^{512}$ wieder eine Divisionsalgebra definiert werden? Adolf Hurwitz konnte 1898 zeigen, dass dies zumindest für euklidisch normierte Divisionsalgebren, welche die Produktregel $|x \cdot y| = |x| \cdot |y|$ erfüllen, nicht möglich ist.

Satz. Sei $\mathbb{R}^n$ eine endlichdimensionale, euklidisch normierte Divisionsalgebra mit $|x \cdot y| = |x| \cdot |y|$ für alle $x, y \in \mathbb{R}^n$. Dann ist $\mathbb{R}^n$ isomorph zu $\mathbb{R}$, $\mathbb{C}$, $\mathbb{H}$ oder $\mathbb{O}$.

Beweisidee. Es lässt sich zeigen, dass jene Divisionsalgebra alternativ sein muss [20, S. 223–224] und ein Einselement besitzt [20, S. 186]. Die Aussage folgt durch Anwenden von Satz 7.3.2. □

Ein weiteres Resultat geht auf Michel Kervaire und John Milnor zurück. Sie bewiesen, dass man auch ohne der Forderung nach Normiertheit Divisionsalgebren nur in den Dimensionen 1, 2, 4 und 8 finden kann [20, S. 233].

Ist daher mit dieser Arbeit das Feld der Zahlenbereichserweiterungen vollständig abgesteckt? Keineswegs! Wir haben uns lediglich den bekannteren Zahlenbereichen gewidmet. Stellen wir bei den Erweiterungen andere Probleme in den Vordergrund, oder modifizieren verwendete Begriffe, gelangen wir zu gänzlich neuen Zahlenbereichen.

Ändern wir etwa die Metrik $|\cdot|$ auf den rationalen Zahlen auf die p -adische Metrik (wobei p für eine beliebige Primzahl steht), erhalten wir die p -adischen Zahlen als vollständige Körpererweiterung der rationalen Zahlen. Die p -adischen Zahlen werden in der Zahlentheorie verwendet, um dieser analytische Methoden zur Verfügung zu stellen.

Leibniz und Newton verwendeten für die von ihnen entwickelte Differentialrechnung das Konzept unendlich kleiner (infinitesimaler) Zahlen, die in unserem Standardmodell der reellen Zahlen nicht existieren können. Die hyperreellen Zahlen erweitern die reellen Zahlen um jene infiniten und infinitesimalen Zahlen, wodurch das Feld der Nonstandardanalysis eröffnet wird. Durch hyperreelle Zahlen ist es möglich, zum Begriff

der Ableitung und des Integrals zu kommen, ohne auf Grenzprozesse zurückgreifen zu müssen.

Kardinalzahlen wurden bereits in Abschnitt 4.7 angesprochen, und dort zur Beschreibung der Mächtigkeiten der bekannten Zahlenmengen verwendet. Auf den Kardinalzahlen, welche die natürlichen Zahlen umfassen, lassen sich wieder Rechenoperationen einführen, und deren Eigenschaften studieren.

Neben der Konstruktion neuer Zahlenbereiche bietet auch die Untersuchung von Teilstrukturen bereits konstruierter Zahlenmengen interessante Möglichkeiten. So bilden etwa die mit Zirkel und Lineal konstruierbaren Zahlen einen Teilkörper der reellen Zahlen. Durch den Fokus auf die algebraische Struktur hinter solchen Konstruktionen konnte die jahrtausendealte Frage um die Möglichkeit der Würfelverdoppelung und der Winkeldreiteilung beantwortet werden.

Zusammenfassung

Diese Diplomarbeit beschäftigt sich mit der mathematisch sauberen Konstruktion und den Eigenschaften von Zahlenbereichen. Sie beginnt mit den Axiomen der Mengenlehre nach Zermelo und Fraenkel (ZFC) und erweitert die daraus konstruierten natürlichen Zahlen auf die ganzen, rationalen, reellen und komplexen Zahlen, bis hin zu den Quaternionen und Oktonionen.

Auf allen genannten Zahlenbereichen werden die bekannten Operationen eingeführt und auf strukturalgebraische Eigenschaften untersucht. Neben einem historischen Abriss über die Gründe, welche zu den Erweiterungen geführt haben, werden auch zentrale Eigenschaften der Zahlenbereiche angesprochen und bewiesen, etwa die Wohlordbarkeit der natürlichen Zahlen oder die Vollständigkeit der reellen Zahlen.

Abstract

This thesis deals with constructing the number systems and showing their properties. The natural numbers are constructed from the axioms of set theory by Zermelo and Fraenkel (ZFC) and subsequently extended to the integers, the rational, real and complex numbers, up to the quaternions and octonions.

We introduce the well-known operations on all those number systems, and examine them for their algebraic properties. An outline of the historical reasons for extending the notion of numbers aside, we present and prove important properties of the number systems, for example the well-orderability of the natural numbers or the completeness of the real numbers.

Abbildungsverzeichnis

1.1	Gruppierte Strichliste	6
1.2	Von Homo erectus gekerbter Knochen	6
1.3	Bijektive Abbildung	8
2.1	Japanische Zeichnung eines Rechenbretts	36
2.2	Gitterdarstellung der Äquivalenzklassen in den ganzen Zahlen	39
3.1	Gitterdarstellung der Äquivalenzklassen in den rationalen Zahlen . .	64
3.2	$\mathbb{Q}$ als kleinster Körper über $\mathbb{Z}$	80
4.1	Regelmäßiges Fünfeck mit Diagonalen	85
4.2	Cantors erstes Diagonalargument	114
5.1	Polardarstellung einer komplexen Zahl	128
5.2	Komplexe Lösungen von $x^5 = 1$	129
6.1	Grafische Multiplikationstabelle in $\mathbb{H}$	139

Literatur

- [1] Martin Aigner und Günter Ziegler. *Das BUCH der Beweise*. 3. Aufl. Springer, 2010.
- [2] John Baez. „The Octonions“. In: *Bulletin of the American Mathematical Society* 39.2 (2002), S. 145–205.
- [3] Martin Barner und Friedrich Flohr. *Analysis 1*. 4. Aufl. Walter de Gruyter, 1991.
- [4] Andreas Büchter. „Das Spiralprinzip“. In: *mathematik lehren* (Feb. 2014), S. 2–9.
- [5] David Burton. *The History of Mathematics: An Introduction*. 7. Aufl. Mc Graw Hill, 2011.
- [6] Chris Caldwell und Yeng Xiong. „What is the smallest prime?“ In: *Journal of Integer Sequences* 15.2 (2012), S. 3.
- [7] Georg Cantor. „Beiträge zur Begründung der transfiniten Mengenlehre“. In: *Mathematische Annalen* 46.4 (1895), S. 481–512.
- [8] Georg Cantor. *Briefe*. Hrsg. von Herbert Meschkowski. Springer, 1991.
- [9] Georg Cantor. „Ein Beitrag zur Mannigfaltigkeitslehre“. In: *Journal für die reine und angewandte Mathematik* 84 (1878), S. 242–258.
- [10] Susan Carey. *The Origin of Concepts*. Oxford University Press, 2009.
- [11] James R Choike. „The pentagram and the discovery of an irrational number“. In: *The Two-Year College Mathematics Journal* 11.5 (1980), S. 312–316.
- [12] Johann Cigler. *Körper - Ringe - Gleichungen*. 1995. URL: <https://homepage.univie.ac.at/johann.cigler/preprints/algebra.pdf> (besucht am 20.08.2018).
- [13] Wikipedia contributors. *Bijektive Funktion*. Bildquelle. 2004. URL: https://de.wikipedia.org/wiki/Bijektive_Funktion (besucht am 13.02.2018).
- [14] Wikipedia contributors. *Fundplatz Bilzingsleben*. Bildquelle. 2004. URL: https://de.wikipedia.org/wiki/Fundplatz_Bilzingsleben (besucht am 10.02.2018).
- [15] John Conway und Derek Smith. *On Quaternions and Octonions*. A K Peters, 2003.

- [16] Rainer Danckwerts und Dankwart Vogel. „Vollständigkeit und Irrationalität – ein schwieriges Geschäft“. In: *Praxis der Mathematik in der Schule* (Okt. 2006), S. 29–32.
- [17] Tobias Dantzig. *Number: The language of science*. Penguin, 2007.
- [18] Richard Dedekind. *Was sind und was sollen die Zahlen? Stetigkeit und irrationale Zahlen*. Mit Erläut. von Stefan Müller-Stach. Springer, 2017.
- [19] Oliver Deiser. *Einführung in die Mengenlehre*. 2. Aufl. Springer, 2004.
- [20] Heinz-Dieter Ebbinghaus u. a. *Zahlen*. 3. Aufl. Springer, 1992.
- [21] Douglas Farenick. *Algebras of Linear Transformations*. Springer, 2001.
- [22] Graham Flegg. *Numbers. Their History and Meaning*. Penguin, 1984.
- [23] Adolf Fraenkel. „Zu den Grundlagen der Cantor-Zermeloschen Mengenlehre“. In: *Mathematische Annalen* 86.3-4 (1922), S. 230–237.
- [24] Kurt Gödel. „Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme I“. In: *Monatshefte für Mathematik und Physik* 38.1 (1931), S. 173–198.
- [25] Gilbert Greefrath u. a. *Didaktik der Analysis*. Springer, 2016.
- [26] Andrew Hanson. *Visualizing Quaternions*. Morgan Kaufmann, 2006.
- [27] Lisa Hefendehl-Hebeker und Susanne Prediger. „Unzählig viele Zahlen: Zahlbereiche erweitern – Zahlvorstellungen wandeln“. In: *Praxis der Mathematik in der Schule* (Okt. 2006), S. 1–7.
- [28] David Hilbert. „Mathematische Probleme“. In: *Nachrichten von der Gesellschaft der Wissenschaften zu Göttingen* 3 (1900), S. 253–297.
- [29] Luke Hodgkin. *A History of Mathematics: From Mesopotamia to Modernity*. Oxford University Press, 2005.
- [30] Adolf Hurwitz. „Ueber die Composition der quadratischen Formen von beliebig vielen Variablen“. In: *Nachrichten von der Gesellschaft der Wissenschaften zu Göttingen* (1898), S. 309–316.
- [31] Georges Ifrah. *Universalgeschichte der Zahlen*. 2. Aufl. Campus Verlag, 1991.
- [32] K. Imaeda und Mari Imaeda. „Sedenions: algebra and analysis“. In: *Applied Mathematics and Computation* 115.2-3 (2000), S. 77–88.
- [33] Asaf Karagila. *There exists an injection from X to Y if and only if there exists a surjection from Y to X* . 2012. URL: <https://math.stackexchange.com/q/192486> (besucht am 24. 09. 2018).
- [34] Günther Karigl. *Zahlenbereichserweiterungen*. Vorlesungsskriptum. 2017.

-
- [35] Jürg Kramer und Anna-Maria von Pippich. *Von den natürlichen Zahlen zu den Quaternionen*. Springer, 2013.
- [36] Bernhard Krön. *Einführung in die Analysis*. Vorlesungsskriptum. 2013.
- [37] Günther Malle. „Zahlen fallen nicht vom Himmel“. In: *mathematik lehren* (Mai 2007), S. 4–11.
- [38] Kenneth O. May. „The Impossibility of a Division Algebra of Vectors in Three Dimensional Space“. In: *The American Mathematical Monthly* 73.3 (1966), S. 289–291.
- [39] Paul J. Nahin. *An Imaginary Tale. The Story of $\sqrt{-1}$* . Princeton University Press, 1998.
- [40] Joseph Needham. *Science and Civilisation in China*. Bd. 3. Cambridge University Press, 1959.
- [41] Adolf Leo Oppenheim. „On an operational device in Mesopotamian bureaucracy“. In: *Journal of Near Eastern Studies* 18.2 (1959). Zitiert in [31], S. 121–128.
- [42] Republik Österreich. *Lehrplan der Volksschule*. URL: https://bildung.bmbwf.gv.at/schulen/unterricht/lp/lp_vs_gesamt_14055.pdf (besucht am 18.05.2018).
- [43] Republik Österreich. *Lehrplan Mathematik AHS Oberstufe*. URL: <https://www.ris.bka.gv.at/GeltendeFassung.wxe?Abfrage=Bundesnormen&Gesetzesnummer=10008568&FassungVom=2018-09-01> (besucht am 26.11.2018).
- [44] Republik Österreich. *Lehrplan Mathematik AHS Unterstufe*. URL: https://bildung.bmbwf.gv.at/schulen/unterricht/lp/ahs14_789.pdf (besucht am 18.05.2018).
- [45] Friedhelm Padberg. *Didaktik der Arithmetik*. 2. Aufl. BI Wissenschaftsverlag, 1992.
- [46] Richard Palais. „The classification of real division algebras“. In: *The American Mathematical Monthly* 75.4 (1968), S. 366–368.
- [47] Giuseppe Peano. *Arithmetices principia, nova methodo exposita*. 1889.
- [48] Kristina Reiss und Gerald Schmieder. *Basiswissen Zahlentheorie*. Springer, 2005.
- [49] Lucio Russo. *Die vergessene Revolution oder die Wiedergeburt des antiken Wissens*. Springer, 2005.
- [50] Hermann Schichl und Roland Steinbauer. *Einführung in das mathematische Arbeiten*. 2. Aufl. Springer, 2012.
- [51] Wolfgang Schneider und Ulman Lindenberger. *Entwicklungspsychologie*. 7. Aufl. Beltz, 2012.
- [52] Silke Thies. „Komplexe Zahlen“. In: *mathematik lehren* (Apr. 1998), S. 57–61.

- [53] Kurt Von Fritz. „The discovery of incommensurability by Hippasus of Metapontum“. In: *Annals of mathematics* (1945), S. 242–264.
- [54] Bartel Leendert van der Waerden. „Zenon und die Grundlagenkrise der griechischen Mathematik“. In: *Mathematische Annalen* 117.1 (1940), S. 141–161.
- [55] Joseph Patrick Ward. *Quaternions and Cayley Numbers*. Kluwer Academic Publishers, 1997.
- [56] Hans Wußing. *6000 Jahre Mathematik. Eine kulturgeschichtliche Zeitreise - 1. Von den Anfängen bis Leibniz und Newton*. Springer, 2008.
- [57] Ernst Zermelo. „Untersuchungen über die Grundlagen der Mengenlehre. I“. In: *Mathematische Annalen* 65.2 (1908), S. 261–281.
- [58] Seijutsu Sangaku Zue. *Chinesisch-Japanische Rechenstäbchen*. Bildquelle. 1795. URL: https://commons.wikimedia.org/wiki/File:Counting_board.jpg (besucht am 27. 11. 2018).