



universität
wien

BACHELORARBEIT

Titel der Bachelorarbeit

Das Auswahlaxiom, die Kontinuumshypothese und das
Determiniertheitsaxiom

Verfasser

Benjamin Daniel Kattnig

angestrebter akademischer Grad

Bachelor of Science (BSc)

Wien, Februar 2025

Studienkennzahl lt. Studienblatt:

Studienrichtung lt. Studienblatt:

Betreuerin:

UA 033 621

Bachelorstudium Mathematik UG2002

Assoz. Prof. Dr. Sandra Müller

Zusammenfassung

In dieser Arbeit beschäftigen wir uns mit drei Aussagen der Mengenlehre: Dem Auswahlaxiom, der Kontinuumshypothese und dem Determiniertheitsaxiom.

In den ersten beiden Kapiteln werden nach einer Einleitung in das Thema alle benötigten Begriffe eingeführt und zugrunde liegende Sätze bewiesen.

Im Anschluss folgen die Beweise des Cantor-Bernstein-Schröder-Theorems, des Satzes von Hartogs und des Satzes von Sierpiński, welches einen Zusammenhang zwischen einer Form der verallgemeinerten Kontinuumshypothese und dem Auswahlaxiom herstellt.

Zuletzt beschäftigen wir uns mit dem Determiniertheitsaxiom und zeigen, dass es mit dem Auswahlaxiom nicht verträglich ist (nach einem Satz von Gale und Stewart), aber einer Form der Kontinuumshypothese beweist (Satz von Davis).

Inhaltsverzeichnis

1	Einleitung	1
1.1	Mathematische Logik	1
1.2	Mengenlehre	2
1.3	Unabhängigkeit von Aussagen	3
2	Mathematischer Hintergrund	5
2.1	Die Zermelo-Fraenkelsche Mengenlehre ZFC	5
2.2	Ordinalzahlen	7
2.3	Natürliche Zahlen	13
2.4	Mächtigkeit	14
2.5	Kardinalzahlen	15
2.6	Spiele	17
2.7	Topologie	18
3	Resultate der Mengenlehre	20
3.1	Das Cantor-Bernstein-Schröder-Theorem	20
3.2	Der Satz von Hartogs	21
3.3	Der Satz von Sierpiński	22
3.4	Ein Satz von Gale und Stewart	27
3.5	Der Satz von Davis	27
4	Fazit	33
	Literatur	34

1 Einleitung

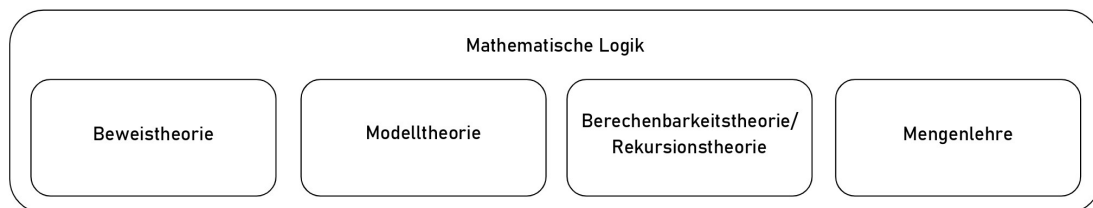
Bevor wir inhaltlich in das Thema dieser Arbeit einsteigen, möchte ich meiner Betreuerin Assoz. Prof. Dr. Sandra Müller sowie der Doktorandin Lena Wallner MSc, die mich beide stets hilfsbereit unterstützt haben, meinen herzlichen Dank aussprechen.

Wir beschäftigen uns in dieser Arbeit mit einigen Resultaten der Mengenlehre, einem Teilgebiet der mathematischen Logik, das heute axiomatisches Fundament vieler Teile der Mathematik ist. Deswegen wollen wir zunächst kurz darauf eingehen, was wir unter mathematischer Logik und naiver bzw. axiomatischer Mengenlehre verstehen. Vergleiche für dieses Kapitel insbesondere [1], [2] und [3].

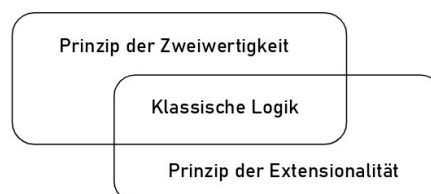
1.1 Mathematische Logik

Logik bezeichnet in der Alltagssprache häufig vor allem rationales Denken. Im wissenschaftlichen Bereich ist der Begriff deutlich enger gefasst und beschäftigt sich mit der abstrakten Analyse von Schlüssen. Dabei spielt er vor allem in der Philosophie, der theoretischen Informatik und der Mathematik eine Rolle. Wir wollen uns ausschließlich mit letzterer beschäftigen, betrachten also mathematische Logik.

Diese lässt sich sowohl inhaltlich als auch strukturell weiter unterteilen: Inhaltlich in vier Teilgebiete, nämlich Beweistheorie, Modelltheorie, Berechenbarkeitstheorie/Rekursionstheorie und Mengenlehre – wir legen den Fokus auf letztere.



Strukturell in verschiedene Eigenschaften wie die Anzahl an möglichen Wahrheitswerten oder welche Informationen genügen, um den Wahrheitswert einer Aussage zu bestimmen. In unserem Framework beschäftigen wir uns ausschließlich mit *klassischer Logik*. In der Literatur ist der Begriff der klassischen Logik leider nicht eindeutig definiert; für uns aber bedeutet er, dass jeder Aussage genau einer von zwei Wahrheitswerten zugeordnet wird (*Prinzip der Zweiwertigkeit*) und der Wahrheitsgehalt einer Aussage eindeutig durch die Wahrheitswerte ihrer Teilaussagen bestimmt ist (*Prinzip der Extensionalität*).



Natürlich gibt es neben der klassischen Logik andere Logiken, die mindestens eines dieser Prinzipien verletzen – z.B. mehrwertige Logiken oder intuitionistische Logik.

Wichtig: Das Prinzip der Extensionalität gilt es nicht mit dem Extensionalitätsaxiom, das wir später kennenlernen werden, zu verwechseln. Während beide den gleichen Sinn haben – nämlich, dass ein Objekt durch seinen Umfang, d.h. seine Extension, bereits genau bestimmt ist – beziehen sie sich auf unterschiedliche Objekte: Das Prinzip bezieht sich auf Aussagen, d.h. der Umfang besteht aus Teilaussagen, während das Axiom sich auf Mengen bezieht, der Umfang also aus Elementen besteht.

Mathematische Logik gibt uns mit der Beweistheorie den entscheidenden Rahmen, um logische formale Systeme (auch *Kalküle* genannt) zu definieren. Hierbei handelt es sich um eine Menge an wohldefinierten Aussagen, die in der Regel anhand eines Alphabets und festen Bildungsregeln entstehen und daher zumeist rekursiv aufzählbar sind, sowie einer Teilmenge an *Axiomen*, die als Basis aller beweisbaren Aussagen dienen. Dazu kommen *Schlussregeln*, die erlauben, aus Axiomen oder bereits bewiesenen Aussagen auf neue Aussagen zu schließen. Das prominenteste Beispiel für eine Schlussregel ist der *Modus Ponens*, der aus den Aussagen A und $A \rightarrow B$ die Folgerung B erlaubt. Definiert man nun einen Beweis als ein Tupel an Aussagen, von denen jede entweder ein Axiom ist oder anhand einer Schlussregel aus vorherigen Aussagen entstanden ist, so hat man ein geeignetes System, um rein formal Aussagen beweisen bzw. treffen-der ableiten zu können. Dieser gesamte Prozess läuft rein *syntaktisch* ab, das bedeutet, es werden ausschließlich „tote“ Zeichenketten betrachtet, die anhand fester Regeln manipuliert werden, ohne, dass eine Bedeutung oder der Begriff der Wahrheit eine Rolle gespielt hätte. Letztere treten erst bei einer *semantischen* Betrachtung eine Rolle, d.h. in der Analyse von Modellen.

1.2 Mengenlehre

Mengenlehre ist (unter anderem) deswegen für die gesamte Mathematik so interessant, weil sich weite Teile der Mathematik auf den Begriff der Menge zurückführen lassen. So kann beispielsweise die natürliche Zahl 0 auf die leere Menge $\emptyset$ zurückgeführt werden und alle größeren natürlichen Zahlen n rekursiv als Mengen der bereits definierten Zahlen 0 bis $n - 1$ definiert werden, d.h. $n = \{0, \dots, n - 1\}$. Ein anderes Beispiel wäre der Begriff der Funktion – jede Funktion f können wir als Menge aller Paare (x, y) auffassen, von denen x unter der Funktion f dem Wert y zugeordnet wird, und der Begriff eines Paares (a, b) wiederum lässt sich z.B. über die Definition $\{a, \{a, b\}\}$ als Menge auffassen. Somit wird aus einer Funktion nichts anderes als eine Menge von Mengen.

Daher ist die Frage, *wie* wir Mengen definieren, von großer Bedeutung. Der *naive* Weg besteht darin, zu jeder möglichen Eigenschaft die Existenz der Menge aller Objekte, die über die gewählte Eigenschaft verfügen, zu postulieren. So lassen sich schnell und einfach alle benötigten Mengen definieren. Doch dieser Ansatz hat ein entscheidendes Problem: Er ist widersprüchlich. Genauer gesagt können wir das *Russel'sche Paradoxon* betrachten: So führt die Existenz der Menge a aller Mengen, die sich nicht selbst enthalten (d.h. die zugehörige Eigenschaft ist $b \notin b$), auf den Widerspruch $a \in a$ und $a \notin a$.

Somit müssen wir einen komplexeren Ansatz wählen, der in der *axiomatischen* Mengenlehre liegt. Hier verwenden wir genau die weiter oben beschriebene Definition eines

formalen Systems. Im nächsten Kapitel werden wir mit dem System $ZF(C)$ eine konkrete, axiomatische Mengenlehre kennenlernen. Dabei steht die Abkürzung ZFC für Zermelo-Fraenkelsche Mengenlehre und ZF für die gleiche, nur um das Auswahlaxiom reduzierte Mengenlehre (C steht für *choice*).

1.3 Unabhängigkeit von Aussagen

Angenommen, wir haben ein formales System definiert, einige Aussagen bewiesen und mithilfe eines Modells aus der Modelltheorie eine Vorstellung davon, welche Aussagen wahr und welche falsch sind. Es stellt sich nun die fundamentale Frage, ob alle wahren Aussagen beweisbar sind.

Als Beispiel wollen wir die Gruppentheorie betrachten, die durch die folgenden drei Axiome begründet wird:

- $\forall x \forall y \forall z ((x \circ y) \circ z) = (x \circ (y \circ z))$
- $\exists e \forall x e \circ x = x$
- $\forall x \exists y x \circ y = e$

Ein mögliches Modell dafür ist die Klein'sche Vierergruppe, oder die Gruppe mit nur einem Element, denn beide erfüllen alle Axiome. Doch was ist mit der Aussage $\exists x \forall y x = y$ (in Worten: es gibt nur ein Element)? Offensichtlich ist sie im Modell der Klein'schen Vierergruppe falsch. Daher kann sie (vorausgesetzt die Axiome sind widerspruchsfrei) nicht aus den Axiomen abgeleitet werden. Andererseits ist sie in der einelementigen Gruppe wahr. Daher kann ihre Negation nicht aus den Axiomen abgeleitet werden, denn sonst könnte die einelementige Gruppe ja kein Modell sein. Insgesamt lässt sich weder die Aussage noch ihre Negation ableiten – wir sagen, dass sie *unabhängig* von den Axiomen der Gruppentheorie ist, die Axiome sind also nicht stark genug, um diese Frage zu beantworten. Damit ist die Gruppentheorie *unvollständig*. Während das bei der Gruppentheorie kein Problem ist, weil wir an unterschiedlichen Gruppen interessiert sind, so sind solche „offenen Fragen“ in anderen Bereichen, z.B. in der Zahlentheorie oder eben der Mengenlehre, unerwünscht. Die Aussage des sogenannten *ersten Gödel'schen Unvollständigkeitssatz* ist gerade, dass jedes widerspruchsfreie rekursiv aufzählbare formale System, das mächtig genug ist, um darin Arithmetik betreiben zu können (wie es die Zahlentheorie offensichtlich ist und die Mengenlehre, wie wir weiter oben erläutert haben, erst recht), unvollständig ist. Aussagen dieser Art treten hier also notwendigerweise auf. Darüber hinaus – dies ist die Aussage des *zweiten Gödel'schen Unvollständigkeitssatzes* – kann jedes solche System seine eigene Widerspruchsfreiheit nicht beweisen. Daher bleibt die Widerspruchsfreiheit von ZF bzw. ZFC nur eine Annahme, die wir oft stillschweigend treffen.

Betrachten wir nun die Aussagen des Titels dieser Arbeit: Das Auswahlaxiom, die (verallgemeinerte) Kontinuumshypothese und das Determiniertheitsaxiom. Wir wollen die Begriffe erst im nächsten Kapitel in ihrem jeweiligen Kontext einführen, aber uns jetzt schon fragen, ob sie beweisbar, widerlegbar oder unabhängig von entsprechenden

Axiomensystemen sind. Die Antwort auf diese Frage lautet: Sofern ZF bzw. ZFC widerspruchsfrei ist, ist das Auswahlaxiom von ZF und sowohl die verallgemeinerte als auch die spezielle Kontinuumshypothese von ZFC (und damit insbesondere auch von ZF) unabhängig. Dennoch herrscht ein enger Zusammenhang zwischen beiden: So impliziert die verallgemeinerte Kontinuumshypothese in der Potenzmengenform (gemeinsam mit den Axiomen von ZF) das Auswahlaxiom. Genau diesen Satz, bekannt als Satz von Sierpiński, beweisen wir im Laufe dieser Arbeit (Satz 3.4).

Das Determiniertheitsaxiom hingegen ist ein komplizierterer Fall. Aufgrund eines Satz von Gale und Stewart, den wir ebenso beweisen werden (Satz 3.9), impliziert es in ZF die Negation des Auswahlaxioms. Daher kann weder ZF noch ZFC das Determiniertheitsaxiom beweisen (sonst wäre das Auswahlaxiom nicht unabhängig von ZF bzw. ZFC nicht widerspruchsfrei). Die andere Richtung, d.h. ob die Negation des Determiniertheitsaxioms beweisbar ist, hängt vom entsprechenden System ab: In ZFC ist sie offensichtlich beweisbar (dies ist direkt aus dem eben erwähnten Satzes von Gale-Stewart ersichtlich), während die Frage in ZF ungelöst ist. Weder gibt es einen Beweis dafür, dass ZF die Negation des Determiniertheitsaxioms impliziert, noch gibt es einen Beweis dafür, dass es einen solchen Beweis nicht geben könne, sofern man nur die Widerspruchsfreiheit von ZF voraussetzt. Denn für letzteres müsste man ein Modell konstruieren, in dem das Determiniertheitsaxiom wahr ist – dies ist zwar möglich, aber nur mit wesentlich stärkeren Annahmen als nur der Konsistenz von ZF.

2 Mathematischer Hintergrund

Vergleiche für die ersten fünf Abschnitte dieses Kapitels *Set Theory and the Continuum Problem* von Raymond M. Smullyan und Melvin Fetting ([4]), sowie [1] und [5] und für die Axiome darüber hinaus auch [6].

Wir beginnen nun mit axiomatischer Mengenlehre. Das Fundament, das wir hierfür verwenden, ist die Prädikatenlogik erster Stufe mit Gleichheit. Eine Definition dieser kann z.B. in [2] nachgelesen werden. Hat man ein Kalkül der Prädikatenlogik erster Stufe aufgebaut, so fügt man zunächst ein zweistelliges Prädikat zur Sprache hinzu, nämlich die *Elementrelation* $\in$. Anschließend definiert man einige *Axiome* als Grundlage. Die Sammlung aller Aussagen, die sich aus diesen Axiomen mithilfe der logischen Axiome und Schlussregeln des gewählten Kalküls ableiten lassen, heißt *Theorie* dieser Axiome.

Zunächst führen wir die Abkürzungen

$\forall x \in y \phi(x)$ für $\forall x (x \in y \rightarrow \phi(x))$ sowie

$\exists x \in y \phi(x)$ für $\exists x (x \in y \wedge \phi(x))$ für jede Formel ϕ und

$\exists! x \phi(x)$ für $\exists x (\phi(x) \wedge \forall y (\phi(y) \rightarrow y = x))$ für jede Formel ϕ und darüber hinaus

$x \subseteq y$ (sprich: x ist eine *Teilmenge* von y) für $\forall a \in x a \in y$

ein.

2.1 Die Zermelo-Fraenkelsche Mengenlehre ZFC

Definition 2.1 (Axiome von Zermelo-Fraenkel ohne Auswahl).

Das formale System ZF ist diese Liste an Axiomen.

- *Extensionalitätsaxiom.* Zwei Mengen sind gleich, wenn sie dieselben Elemente haben: $\forall x \forall y (\forall a (a \in x \leftrightarrow a \in y) \rightarrow x = y)$.¹
- *Leermengenaxiom.* Die leere Menge $\emptyset$ existiert: $\exists x \forall y y \notin x$.
- *Paarungsaxiom.* Für alle Mengen x und y existierte die Menge $z = \{x, y\}$, die genau x und y als Elemente hat: $\forall x \forall y \exists z \forall a (a \in z \leftrightarrow (a = x \vee a = y))$.

Das Paarungsaxiom sichert insbesondere die Existenz von $\{x\}$ für jede Menge x .

- *Vereinigungsaxiom.* Für jede Menge x existiert die Menge $y = \bigcup x$, die genau aus allen Elementen aller Elemente von x besteht: $\forall x \exists y \forall a (a \in y \leftrightarrow (\exists z \in x a \in z))$.

Gemeinsam mit dem Paarungsaxiom impliziert das Vereinigungsaxiom insbesondere die Existenz von $x \cup y$ für zwei Mengen x und y .

¹Wir arbeiten in Prädikatenlogik 1. Stufe mit Gleichheit. Würden wir auf diese semantische Definition der Gleichheit verzichten und stattdessen Gleichheit als Prädikat – ähnlich der Elementrelation – einführen, so würden wir das Extensionalitätsaxiom dahingehend abändern, dass zwei Mengen *genau dann* gleich sind, wenn sie dieselben Elemente haben, und hinzufügen, dass gleiche Mengen *Elemente derselben Mengen* sind. In unserem Fall ergeben sich beide Eigenschaften direkt durch die semantische Definition der Gleichheit, weshalb wir sie nicht im Extensionalitätsaxiom enthalten haben.

- *Aussonderungsaxiomenschema.* Zu jeder Menge x und jeder prädikatenlogischen Formel ϕ mit Parametern $p_1, \dots, p_n$ existiert die Menge y aller Elemente aus x , die ϕ erfüllen: $\forall x \forall p_1 \dots \forall p_n \exists y \forall a (a \in y \leftrightarrow (a \in x \wedge \phi(a, p_1, \dots, p_n)))$.

Das Aussonderungsaxiomenschema liefert (wie auch das gleich folgende Ersetzungsaxiomenschema) ein einzelnes Axiom für jede prädikatenlogische Formel ϕ , weshalb ZF über unendlich viele Axiome verfügt. Durch ein passendes Aussonderungsaxiom können wir für jede Formel ϕ und jede Menge x die Menge $y = \{a \in x \mid \phi(a)\}$ bilden. Insbesondere lässt sich damit $x \cap y$ durch $\{a \in x \mid a \in y\}$ realisieren.

- *Ersetzungsaxiomenschema.* Sei x eine Menge und ϕ eine prädikatenlogische Formel mit Parametern $p_1, \dots, p_n$, die wir als eine Funktion über den Variablen a und b auffassen können. Dann ist das Bild von x unter ϕ eine Menge:

$$\forall p_1 \dots \forall p_n \forall a \exists! b \phi(a, b, p_1, \dots, p_n) \rightarrow$$

$$\forall x \exists y \forall z (z \in y \leftrightarrow \exists v \in x \phi(v, z, p_1, \dots, p_n)).$$

- *Potenzmengenaxiom.* Für jede Menge x existiert die Menge $y = \mathcal{P}(x)$ aller Teilmengen von x : $\forall x \exists y \forall a (a \in y \leftrightarrow a \subseteq x)$.

Wir bezeichnen die Menge $\mathcal{P}(x)$ als *Potenzmenge* von x . Es gilt beispielsweise $\mathcal{P}(\{0, 1\}) = \{\emptyset, \{0\}, \{1\}, \{0, 1\}\}$.

- *Unendlichkeitsaxiom.* Es gibt eine Menge x , die die leere Menge $\emptyset$ und für jede enthaltene Menge y auch $y \cup \{y\}$ enthält: $\exists x (\emptyset \in x \wedge \forall y \in x (y \cup \{y\} \in x))$.

Wir werden diese Menge insbesondere für die Definition der natürlichen Zahlen benötigen.

- *Fundierungsaxiom.* Jede nichtleere Menge x hat ein Element y , das zu x disjunkt ist: $\forall x (x \neq \emptyset \rightarrow \exists y \in x (x \cap y = \emptyset))$.

Das Fundierungsaxiom verhindert die Existenz unendlicher $\in$ -absteigender Folgen ($x_0 \ni x_1 \ni x_2 \ni \dots$): Angenommen, es gäbe eine solche Folge $(x_n)_n$ mit $x_n \ni x_{n+1}$. Dann existiert die Menge $\{x_0, x_1, x_2, \dots\}$ nach einer Instanz des Ersetzungsaxiomenschema. Diese Menge widerspricht jedoch dem Fundierungsaxiom, weil für jedes Element $x_i \in \{x_0, x_1, x_2, \dots\}$ der Schnitt $x_i \cap \{x_0, x_1, x_2, \dots\}$ die Menge x_{i+1} enthält, also nichtleer ist.

Außerdem sichert das Fundierungsaxiom die Irreflexivität von $\in$ (d.h. keine Menge enthält sich selbst): Angenommen, es gibt eine Menge a mit $a \in a$. Dann widerspricht die Menge $\{a\}$ dem Fundierungsaxiom, weil sie kein Element enthält, dass disjunkt zu $\{a\}$ ist.

Definition 2.2 (Axiome von Zermelo-Fraenkel mit Auswahl).

Das formale System *ZFC* besteht aus den Axiomen von ZF und dem *Auswahlaxiom*, welches für jede Menge x nichtleerer und paarweise disjunkter Mengen die Existenz einer Auswahlmenge y postuliert, die aus jedem Element von x genau ein Element enthält:

$$\forall x ((\forall x_1 \in x \forall x_2 \in x x_1 \neq \emptyset \wedge (x_1 = x_2 \vee x_1 \cap x_2 = \emptyset)) \rightarrow \exists y \forall x_0 \in x \exists! y_0 \in y y_0 \in x_0).$$

Bemerkung. Das Auswahlaxiom wird oft mit AC (aus dem Englischen für *axiom of choice*) abgekürzt. Es ist vom System ZF unabhängig, lässt sich also (vorausgesetzt ZF ist widerspruchsfrei) weder beweisen noch widerlegen.

Die Einführung dieser Axiome erlaubt uns, einige wichtige Begriffe zu definieren. Wo nicht explizit anders ausgedrückt, ist der formale Rahmen für alle in dieser Arbeit folgenden Definitionen und Beweise stets ZF.

Definition 2.3 (Paar, Tupel und kartesisches Produkt). Seien x und y Mengen. Dann definieren wir das *Paar* oder *2-Tupel* von x und y als $(x, y) := \{\{x\}, \{x, y\}\}$ und das *n -Tupel* rekursiv über $(x_1, \dots, x_n) := ((x_1, \dots, x_{n-1}), x_n)$. Das *kartesische Produkt* von x und y ist $x \times y := \{(a, b) \in \mathcal{P}(\mathcal{P}(x \cup y)) \mid a \in x \text{ und } b \in y\}$.

Definition 2.4 (Relation und Funktion). Seien x und y Mengen. Dann heißt jede Teilmenge r von $x \times y$ eine (binäre) *Relation auf x und y* . Falls $x = y$ gilt, so sprechen wir von einer *Relation auf x* . Wir schreiben arb , falls $(a, b) \in r$ gilt.

Eine Relation f auf x und y , die die Eigenschaft $\forall a \in x \exists! b \in y \, afb$ erfüllt, heißt *Funktion von x nach y* . Wir schreiben dann $f: x \rightarrow y$ und $f(a) = b$, falls $(a, b) \in f$ gilt.

Definition 2.5 (Menge an Funktionen). Seien x und y Mengen. Dann definieren wir ${}^x y := \{f \in \mathcal{P}(x \times y) \mid f \text{ ist eine Funktion von } x \text{ nach } y\}$.

Definition 2.6 (Lineare Ordnungsrelation). Sei $<$ eine Relation auf einer Menge x . Dann ist $<$ eine *lineare Ordnungsrelation*, wenn $<$ für alle $a, b, c \in x$ die Eigenschaften der

- Irreflexivität $a \not< a$,
- Transitivität $(a < b \wedge b < c) \rightarrow a < c$ und
- Linearität $a < b \vee a = b \vee b < a$

erfüllt. Wir bezeichnen die lineare Ordnung $<$ auch mit $(x, <)$, um auszudrücken, dass sie die Menge x ordnet.

2.2 Ordinalzahlen

Definition 2.7 (Wohlordnung). Sei $<$ eine lineare Ordnungsrelation auf einer Menge x . Dann ist $(x, <)$ eine *Wohlordnung*, falls jede nichtleere Teilmenge von x ein $<$ -kleinstes Element hat, d.h. es gilt $\forall y ((y \subseteq x \wedge y \neq \emptyset) \rightarrow \exists b \in y (\forall a \in y \, b < a))$.

Definition 2.8 (Ordnungsisomorphie). Seien $(x, <)$ und $(y, <)$ zwei Wohlordnungen. Dann bezeichnen wir eine bijektive Funktion $f: x \rightarrow y$ als *Ordnungsisomorphismus*, wenn für alle $a, b \in x$ die Eigenschaft

$$a < b \leftrightarrow f(a) < f(b)$$

erfüllt ist. Wir nennen die Mengen x und y *ordnungsisomorph* und schreiben $(x, <) \equiv (y, <)$, falls es eine solche Funktion gibt.

Definition 2.9 (Anfangsstück). Sei $(x, <)$ eine Wohlordnung und $a \in x$. Dann bezeichnen wir die Wohlordnung $(x_a, <_a)$ mit

- $x_a := \{b \in x \mid b < a\}$ und
- $<_a := \{(b, c) \in < \mid b, c \in x_a\}$

als a -Anfangsstück von $(x, <)$.

Wir schreiben aus Gründen der Leserlichkeit oft nur $(x_a, <)$ statt $(x_a, <_a)$. Es lässt sich leicht zeigen, dass jedes Anfangsstück einer Wohlordnung wieder eine Wohlordnung ist. Wir wollen den Begriff der Wohlordnung mithilfe des Begriffes der *Ordinalzahl* näher analysieren. So steht jede Ordinalzahl als „Prototyp“ für genau eine mögliche Wohlordnung, und jeder Wohlordnung kann eine eindeutige Ordinalzahl zugeordnet werden, wie wir in Kürze beweisen. Zunächst wollen wir den *Vergleichbarkeitssatz für Wohlordnungen*² beweisen und benötigen dafür ein Lemma.

Lemma 2.10. Sei $(x, <)$ eine Wohlordnung und seien $a, b \in x$. Dann gilt

- (a) $(x_a, <) \not\equiv (x, <)$ und
- (b) $a < b$ genau dann, wenn $(x_a, <)$ ein Anfangsstück von $(x_b, <)$ ist, und
- (c) $a = b$ genau dann, wenn $(x_a, <) \equiv (x_b, <)$ gilt.

Beweis. (a) Angenommen, es gäbe ein $a \in x$ mit $(x_a, <) \equiv (x, <)$. Dann gibt es einen Ordnungsisomorphismus $f: x \rightarrow x_a$. Also ist $f(a) \in x_a$. Daher muss $f(a) < a$ sein, weswegen die Menge $y = \{c \in x \mid f(c) < c\}$ nichtleer ist. Sei y_0 das kleinste Element von y . Es gilt also $f(y_0) < y_0$. Da f ein Ordnungsisomorphismus ist, gilt also auch $f(f(y_0)) < f(y_0)$. Damit ist auch $f(y_0) \in y$. Aber $f(y_0) < y_0$, im Widerspruch zur Minimalität von y_0 .

- (b) Es gelte $a < b$. Dann ist $x_a = \{c \in x \mid c < a\} = \{c \in x_b \mid c < a\}$ und $<_a = \{(c, d) \in < \mid c, d \in a\} = \{(c, d) \in <_b \mid c, d \in a\}$, also ist $(x_a, <)$ ein Anfangsstück von $(x_b, <)$.

Es sei $(x_a, <)$ ein Anfangsstück von $(x_b, <)$. Wäre $b < a$, dann wäre $(x_b, <)$ ein Anfangsstück von $(x_a, <)$, woraus leicht folgt, dass $(x_a, <)$ ein Anfangsstück von $(x_a, <)$ ist und damit ein Widerspruch zu (a). Der gleicher Widerspruch würde aus $a = b$ folgen, es gilt also wie erwünscht $a < b$.

- (c) Es gelte $a = b$. Dann ist $(x_a, <) \equiv (x_b, <)$ klar.

Es gelte $(x_a, <) \equiv (x_b, <)$. Wäre $a \neq b$, so würde entweder $a < b$ oder $b < a$ gelten. Damit folgt aus (b), dass entweder x_a ein Anfangsstück von x_b ist oder andersherum. Dies führt jedoch auf einen Widerspruch zu (a).

□

²Den Vergleichbarkeitssatz für Wohlordnungen gilt es nicht mit dem Vergleichbarkeitssatz für Mengen zu verwechseln: Letzterer behauptet, dass für je zwei Mengen a und b die Eigenschaft $a \preceq b \vee b \preceq a$ erfüllt ist. Dies ist eine wesentlich stärkere Aussagen, die in ZF zum Auswahlaxiom äquivalent ist.

Satz 2.11 (Vergleichbarkeitssatz für Wohlordnungen). *Seien $(x, <)$ und $(y, <)$ zwei Wohlordnungen. Dann gilt*

- $(x, <) \equiv (y, <)$, oder
- es gibt ein $a \in x$ sodass $(x_a, <_a) \equiv (y, <)$ gilt, oder
- es gibt ein $b \in y$ sodass $(x, <) \equiv (y_b, <_b)$ gilt.

Beweis. Wir definieren die Menge $f = \{(a, b) \mid a \in x, b \in y \text{ und } (x_a, <) \equiv (y_b, <)\}$.

Schritt 1: Zunächst zeigen wir, dass f dann eine injektive Funktion ist. Seien dazu (a, b_1) und (a, b_2) aus f . Dann gilt $(x_a, <) \equiv (y_{b_1}, <)$ und $(x_a, <) \equiv (y_{b_2}, <)$; also auch $(y_{b_1}, <) \equiv (y_{b_2}, <)$, daher ist nach Lemma 2.10 $b_1 = b_2$ und f damit eine Funktion. Analog erhält man die Injektivität durch Betrachtung von (a_1, b) und (a_2, b) .

Schritt 2: Sei nun $a \in \text{dom}(f)$. Wir zeigen, dass für alle $b \in x$ mit $b < a$ auch $b \in \text{dom}(f)$ ist: Da $a \in \text{dom}(f)$ ist, gibt es $c \in \text{im}(f)$, sodass $(x_a, <) \equiv (x_c, <)$ gilt. Es lässt sich leicht zeigen, dass die Einschränkung der dies bezeugenden Funktion g auf x_b ebenfalls ein Ordnungsisomorphismus ist, es folgt also $(x_b, <) \equiv (x_{g(b)}, <)$. Damit ist $(b, g(b)) \in f$, also wie erwünscht $b \in \text{dom}(f)$. Analog folgt, dass für jedes $a \in \text{im}(f)$ auch alle $b \in y$ mit $b < a$ in $\text{im}(f)$ sind (man betrachte dafür die Umkehrfunktion der Ordnungsisomorphie).

Schritt 3: Nun zeigen wir, dass $f: \text{dom}(f) \rightarrow \text{im}(f)$ selbst ein Ordnungsisomorphismus ist: Injektivität haben wir bereits gezeigt, und Surjektivität folgt durch die Definition. Seien also $a_1, a_2 \in \text{dom}(f)$ mit $a_1 < a_2$ und $f(a_1) = b_1$, $f(a_2) = b_2$. Dann gilt $(y_{b_1}, <) \equiv (x_{a_1}, <)$. Weil $a_1 < a_2$ ist, ist $(x_{a_1}, <)$ ordnungsisomorph zu einem Anfangsstück von $(x_{a_2}, <)$ (Lemma 2.10). Aber $(x_{a_2}, <) \equiv (y_{b_2}, <)$, was insgesamt ergibt, dass $(y_{b_1}, <)$ ordnungsisomorph zu einem Anfangsstück von $(y_{b_2}, <)$ ist. Daher gilt (erneut nach Lemma 2.10) $b_1 < b_2$, also $f(a_1) < f(a_2)$.

Schritt 4: Zuletzt halten wir fest, dass $\text{dom}(f) = x$ oder $\text{im}(f) = y$ gilt: Angenommen, dies gilt nicht. Sei dann a das kleinste Element aus $x \setminus \text{dom}(f)$ und b das kleinste Element aus $y \setminus \text{im}(f)$. Durch Schritt 2 folgt dann $\text{dom}(f) = x_a$ und $\text{im}(f) = y_b$. Wegen Schritt 3 gilt weiter $(x_a, <) \equiv (y_b, <)$. Also ist (a, b) nach Definition ein Element von f , ein Widerspruch.

Im Falle $\text{dom}(f) = x$ und $\text{im}(f) = y$ erhalten wir direkt die Ordnungsisomorphie von x und y . Falls $\text{dom}(f) = x$ ist, aber $\text{im}(f) \neq y$, so ist $(x, <)$ ordnungsisomorph zu einem Anfangsstück von $(y, <)$, und im Falle $\text{dom}(f) \neq x$ und $\text{im}(f) = y$ ist $(y, <)$ ordnungsisomorph zu einem Anfangsstück von $(x, <)$. $\square$

Definition 2.12 (Transitivität). Eine Menge x heißt *transitiv*, wenn $\forall a \in x \ a \subseteq x$ gilt, oder äquivalent dazu $\forall a \ \forall b \ ((a \in b \wedge b \in x) \rightarrow a \in x)$ erfüllt ist.

Definition 2.13 (Ordinalzahl). Eine transitive Menge x , deren Elemente bezüglich $\in$ wohlgeordnet sind, heißt *Ordinalzahl*.

Ordinalzahlen bezeichnen wir in der Regel mit griechischen Kleinbuchstaben.

Definition 2.14 (Ordnung der Ordinalzahlen). Wir definieren durch $\alpha < \beta :\Leftrightarrow \alpha \in \beta$ eine Ordnungsrelation zwischen Ordinalzahlen.

A priori ist nicht klar, wieso $<$ eine Ordnungsrelation ist. Dies erhalten wir erst durch Punkt (c) des folgenden Satzes. Zu beachten ist, dass $<$ keine Ordnungsrelation im Sinne von Definition 2.6 ist, sondern nur auf der Meta-Ebene mithilfe einer prädikatenlogischen Formel definiert ist, weil es die „Menge“ aller Ordinalzahlen, auf der die Relation definiert sein müsste, nicht gibt (Punkt (e)).

Satz 2.15 (Eigenschaften von Ordinalzahlen).

- (a) Jede Ordinalzahl ist die Menge aller kleineren Ordinalzahlen.
- (b) Zwei Ordinalzahlen mit ordnungsisomorphen Wohlordnungen sind gleich.
- (c) Die Ordinalzahlen sind unter $<$ wohlgeordnet.
- (d) Die Vereinigung einer Menge von Ordinalzahlen ist eine Ordinalzahl.
- (e) Die Menge aller Ordinalzahlen existiert nicht.
- (f) Jede wohlgeordnete Menge ist zu genau einer Ordinalzahl ordnungsisomorph.

Beweis.

- (a) Sei α eine Ordinalzahl und $a \in \alpha$. Aufgrund der Transitivität von α ist dann $a \subseteq \alpha$, die Wohlordnung von α lässt sich also auf a einschränken. Darüber hinaus ist a transitiv: Sei $b \in a$ und $c \in b$. Da α als Ordinalzahl transitiv ist, gilt $b \in \alpha$. Damit ist, erneut aufgrund der Transitivität von α , $c \in \alpha$. Da α wohlgeordnet ist und $a \in \alpha$ sowie $c \in \alpha$ ist, gilt $a \in c \vee a = c \vee c \in a$. Angenommen, es gilt $a \in c$. Dann führt $a \in c \in b \in a$ zu einem Widerspruch zum Fundierungsaxiom (betrachte dafür die Menge $\{a, b, c\}$). Auch der Fall $a = c$ widerspricht dem Fundierungsaxiom (betrachte die Menge $\{a, b\}$). Also folgt $c \in a$, womit die Transitivität von a bewiesen ist. Damit ist a eine Ordinalzahl, es gilt also insbesondere $a < \alpha$.

Sei umgekehrt β eine Ordinalzahl mit $\beta < \alpha$. Dann gilt nach Definition direkt $\beta \in \alpha$.

- (b) Seien α und β zwei Ordinalzahlen, sodass $(\alpha, \in) \equiv (\beta, \in)$ gilt. Dann existiert eine bijektive Funktion $f: \alpha \rightarrow \beta$, die die Ordnungsisomorphie bezeugt. Wenn f die Identität ist, so sind wir fertig. Andernfalls gibt es ein minimales $\gamma \in \alpha$ sodass $f(\gamma) \neq \gamma$ gilt. Nach (a) sind γ und $f(\gamma)$ Ordinalzahlen. Damit folgt der Widerspruch $\gamma = \alpha_\gamma = \beta_{f(\gamma)} = f(\gamma)$.
- (c) Seien α , β und γ Ordinalzahlen.

Irreflexivität: Siehe Anmerkung zum Fundierungsaxiom.

Transitivität: Es gelte $\alpha < \beta$ und $\beta < \gamma$. Dann gilt $\alpha \in \beta \in \gamma$, und da γ eine Ordinalzahl ist, ist γ transitiv. Daher folgt wie erwünscht $\alpha \in \gamma$, also $\alpha < \gamma$.

Linearität: Aufgrund des Vergleichbarkeitssatzes von Wohlordnungen (Satz 2.11) tritt einer der drei folgenden Fälle ein: Im ersten Fall gilt $(\alpha, \in) \equiv (\beta, \in)$, woraus durch (b) $\alpha = \beta$ folgt. Im zweiten Fall gibt es ein $a \in \alpha$ sodass $(\alpha_a, \in_a) \equiv (\beta, \in)$ gilt. Dank (a) ist a eine Ordinalzahl und es gilt $(\alpha_a, \in_a) = (a, \in)$. Durch (b) folgt $a = \beta$, d.h. insbesondere $\beta \in \alpha$. Der dritte Fall liefert analog zum zweiten $\alpha \in \beta$.

Wohlordnungseigenschaft: Sei x eine nichtleere Menge an Ordinalzahlen. Angenommen, x hat kein kleinstes Element. Dann gilt für jedes $a \in x$, dass $a \cap x \neq \emptyset$ (andernfalls wäre a das kleinste Element, weil dann für alle $b \in x$ die Eigenschaft $b \notin a$ erfüllt ist, also aufgrund der soeben bewiesenen Linearität die Eigenschaft $b = a \vee a \in b$, d.h. $a \leq b$ gilt). Das ist ein Widerspruch zum Fundierungssaxiom.

- (d) Sei a eine Menge an Ordinalzahlen. Wir zeigen, dass $\bigcup a$ durch $\in$ wohlgeordnet und transitiv ist. Seien $b, c, d \in \bigcup a$. Dann gibt es a_b und a_c in a , sodass $b \in a_b$ und $c \in a_c$.

Irreflexivität: Siehe Anmerkung zum Fundierungssaxiom.

Linearität von $\in$: Da die Ordinalzahlen nach (c) wohlgeordnet sind, gilt $a_b < a_c \vee a_b = a_c \vee a_c < a_b$, also $a_b \in a_c \vee a_b = a_c \vee a_c \in a_b$. In allen Fällen folgt aufgrund der Transitivität von Ordinalzahlen $b, c \in a_b$ oder $b, c \in a_c$ und damit $b \in c \vee b = c \vee c \in b$, weil a_b bzw. a_c eine Ordinalzahl ist.

Transitivität von $\in$: Angenommen, es gilt $b \in c$ und $c \in d$. Aufgrund der soeben bewiesenen Linearität gilt dann $b \in d \vee b = d \vee d \in b$. Die letzteren beiden Fällen führen jedoch zu einem Widerspruch zum Fundierungssaxiom, wenn man die Menge $\{b, c, d\}$ betrachtet. Also gilt $b \in d$.

Transitivität von $\bigcup a$: Es ist $b \subseteq \bigcup a$, weil $b \subseteq a_b$ (Transitivität der Ordinalzahl a_b) und $a_b \subseteq \bigcup a$ gilt.

- (e) Angenommen, die Menge x aller Ordinalzahlen existiert. Dann ist x nach (c) durch $\in$ wohlgeordnet und nach (a) transitiv, also selbst eine Ordinalzahl, woraus $x \in x$ folgt – ein Widerspruch zur Irreflexivität von $\in$.
- (f) Die Eindeutigkeit folgt bereits aus Punkt (b). Angenommen, es gibt eine Wohlordnung $(x, <)$, sodass keine zu $(x, <)$ ordnungsisomorphe Ordinalzahl α existiert. Wir definieren

$$y = \{a \in x \mid \text{es gibt eine Ordinalzahl } \alpha \text{ mit } (x_a, <_a) \equiv (\alpha, \in)\}.$$

Anschließend ordnen wir jedem Element a von y durch die Funktion $f: y \rightarrow im(f)$ die eindeutige Ordinalzahl α zu, die erfüllt, dass $(x_a, <_a) \equiv (\alpha, \in)$ gilt. Dann ist $im(f)$ nach dem Aussonderungssaxiomenschema eine Menge an Ordinalzahlen. Tatsächlich enthält $im(f)$ aber bereits alle Ordinalzahlen:

Angenommen, eine Ordinalzahl β ist nicht in $im(f)$ enthalten. Der Vergleichbarkeitssatz für Wohlordnungen (Satz 2.11) liefert uns die folgenden drei Fälle:

Fall 1: $(\beta, \in) \equiv (x, <)$. Das liefert einen Widerspruch zur ursprünglichen Annahme, dass es für $(x, <)$ keine ordnungsisomorphe Ordinalzahl gäbe.

Fall 2: Ein Anfangsstück von $(x, <)$ ist ordnungsisomorph zu $(\beta, \in)$. Das widerspricht $\beta \notin \text{im}(f)$.

Fall 3: Ein Anfangsstück von $(\beta, \in)$ ist ordnungsisomorph zu $(x, <)$. Es existiert also $\gamma \in \beta$ mit $(\beta_\gamma, \in_\gamma) \equiv (x, <)$. Nach (a) gilt dann $(\gamma, \in) \equiv (x, <)$, im Widerspruch zur ursprünglichen Annahme.

Da sämtliche Fälle in einem Widerspruch münden, müssen alle Ordinalzahlen in β enthalten sein. Die „Menge“ $\text{im}(f)$ aller Ordinalzahlen kann jedoch nach (e) nicht existieren – Widerspruch.

□

Definition 2.16 (Nachfolger einer Ordinalzahl). Sei α eine Ordinalzahl. Dann heißt $\alpha + 1 := \alpha \cup \{\alpha\}$ *Nachfolger* von α .

Es lässt sich leicht beweisen, dass der Nachfolger einer Ordinalzahl α wieder eine Ordinalzahl ist und darüber hinaus die kleinste Ordinalzahl, die größer als α ist. Mithilfe des Begriffes des Nachfolgers einer Ordinalzahl können wir alle Ordinalzahlen in drei Gruppen einteilen: Jede Ordinalzahl α ist entweder gleich $\emptyset$, oder eine *Nachfolgerordinalzahl*, d.h. es gibt eine Ordinalzahl β mit $\beta + 1 = \alpha$, oder weder noch. Im letzten Fall bezeichnen wir α als *Limesordinalzahl*. Mithilfe des *Rekursionssatzes für Ordinalzahlen*, in dem wir alle drei Fälle durchgehen, können wir Operationen bzw. Aussagen für alle Ordinalzahlen definieren bzw. beweisen. Dies erinnert stark an rekursive bzw. induktive Beweise für die Menge der natürlichen Zahlen, nur mit dem Unterschied, dass es drei Fälle statt zwei zu behandeln gibt.

Satz 2.17. *In ZF gilt: Aus der Annahme, jede Menge kann wohlgeordnet werden, folgt das Auswahlaxiom.*

Beweis. Sei x eine Menge disjunkter, nichtleerer Mengen. Nach der Annahme gibt es eine Wohlordnung $<$ für die Menge $\bigcup x$. Wir bilden nun mittels des Aussonderungsaxioms und dieser Wohlordnung die Menge $y = \{a \in \bigcup x \mid \exists u \in x (a \in u \wedge \forall b \in u (b = a \vee a < b))\}$. In Worten beschrieben entspricht y der Menge, die von jedem Element von x das gemäß der Wohlordnung kleinste Element enthält – also genau eines. Damit haben wir eine Auswahlmenge für x gefunden. □

Bemerkung. Die Annahme, jede Menge könne wohlgeordnet werden (oder äquivalent dazu: Für jede Menge x gibt es eine ordnungsisomorphe Ordinalzahl α), ist bekannt als *Wohlordnungssatz*. Es handelt sich dabei um einen *Satz*, weil die Rückrichtung von Satz 2.17 ebenfalls beweisbar ist, der Wohlordnungssatz und das Auswahlaxiom sind also in ZF äquivalent und der Wohlordnungssatz ist damit ein Element der Theorie von ZFC. Wir verzichten auf einen Beweis in die andere Richtung und verweisen z.B. auf [4].

2.3 Natürliche Zahlen

Definition 2.18 (Induktive Menge). Eine Menge x heißt *induktiv*, wenn sie die beiden Eigenschaften

- $\emptyset \in x$ und
- für alle $y \in x$ ist auch $y \cup \{y\} \in x$

erfüllt.

Man beachte, dass $y \cup \{y\}$ aufgrund des Paarungs- und des Vereinigungsaxioms für alle Mengen y existiert.

Definition 2.19 (Menge der natürlichen Zahlen). Die Menge $\omega := \bigcap \{a \mid a \text{ ist induktiv}\}$ heißt *Menge der natürlichen Zahlen*.

Obige Definition können wir nur aufgrund des Unendlichkeitsaxioms verwenden. Dafür bemerken wir zunächst, dass es nach dem Unendlichkeitsaxiom eine induktive Menge x gibt, es gilt also $\bigcap \{a \mid a \text{ ist induktiv}\} \subseteq x$. Daher haben wir

$$\begin{aligned}\omega &= \bigcap \{a \mid a \text{ ist induktiv}\} \\ &= \bigcap \{a \subseteq x \mid a \text{ ist induktiv}\} \\ &= \bigcap \{a \in \mathcal{P}(x) \mid a \text{ ist induktiv}\} \\ &= \{b \in x \mid \forall a \in \mathcal{P}(x) (a \text{ ist induktiv} \rightarrow b \in a)\},\end{aligned}$$

was nach dem Potenzmengenaxiom und einer Instanz des Aussonderungsaxiomenschema eine Menge ist.

Bemerkung. Wir stellen nun fest, dass die auf diese Weise definierten natürlichen Zahlen der bereits bekannten Menge $\mathbb{N}$ entspricht, wenn wir die folgenden Bezeichnungen einführen. Dabei geben wir jedem Element x von ω genau dann den Namen „ n “, wenn es n Elemente hat.

- $0 = \emptyset$
- $1 = \{\emptyset\} = \{0\}$
- $2 = \{\emptyset, \{\emptyset\}\} = \{0, 1\}$
- $3 = \{\emptyset, \{\emptyset\}, \{\{\emptyset, \{\emptyset\}\}\} = \{0, 1, 2\}$

Es lässt sich zeigen, dass jedes Element von ω endlich ist, was bedeutet, dass diese Namenszuordnung wohldefiniert ist. Außerdem ist jede natürliche Zahl eine Ordinalzahl, insbesondere ist die Menge aller natürlichen Zahlen also wohlgeordnet. Daher erhalten wir, dass für je zwei verschiedene natürliche Zahlen m und n entweder $m \in n$ oder $n \in m$ gilt, also ist die Namenszuordnung darüber hinaus injektiv. Und da die Menge $x \cup \{x\}$ genau ein Element mehr als die Menge x hat (wofür wir das Fundierungsaxiom benötigen), gibt es für jedes Element n aus $\mathbb{N}$ auch ein Element aus ω mit n Elementen, die Zuordnung ist also auch surjektiv. Da wir diese Theorie jedoch in dieser Arbeit nicht benötigen, fahren wir ohne Beweise fort.

2.4 Mächtigkeit

Ein der grundlegendsten Fragen in der Mengenlehre ist der Begriff der *Größe* einer Menge, auch *Mächtigkeit* genannt. Im endlichen Fall bereitet der Begriff uns keine Schwierigkeiten, denn wir können die Anzahl an Elementen einer Menge zählen und davon ausgehend Mengen vergleichen. Was aber, wenn eine Menge unendlich groß ist? Sind alle unendlichen Mengen gleich groß, haben also die gleiche Mächtigkeit? Oder gibt es verschiedene Unendlichkeiten? Wie *Georg Cantor* als Erster entdeckte, ist letzteres der Fall. Um das zu beweisen, müssen wir zunächst einmal den Begriff der *Gleichmächtigkeit* zweier Mengen einführen.

Definition 2.20 (Mächtigkeitsrelation). Zwei Mengen x und y heißen *gleichmächtig*, in Zeichen $x \cong y$, wenn es eine Bijektion $f: x \rightarrow y$ gibt. Eine Menge x heißt *weniger oder gleichmächtig* zu einer Menge y , in Zeichen $x \preceq y$, wenn es eine Injektion $f: x \rightarrow y$ gibt. Zuletzt nennen wir eine Menge x *weniger mächtig* als eine Menge y , in Zeichen $x \prec y$, wenn $x \preceq y$ und $x \not\cong y$ gilt.

Satz 2.21 (Alternative Charakterisierung der Mächtigkeitsrelation). *Seien x und y Mengen. Dann gilt $x \preceq y$ genau dann, wenn es eine bijektive Funktion $f: x \rightarrow a$ mit $a \subseteq y$ gibt.*

Beweis. Hinrichtung: Es gelte $x \preceq y$. Dann gibt es eine injektive Funktion $g: x \rightarrow y$. Daher ist die Funktion $h: x \rightarrow \text{im}(g)$ mit $h(z) = g(z)$ für alle $z \in x$ ebenfalls injektiv und per Definition surjektiv, also bijektiv. Ebenso gilt $\text{im}(h) \subseteq y$.

Rückrichtung: Sei $g: x \rightarrow a$ mit $a \subseteq y$ bijektiv. Dann ist die Funktion $h: x \rightarrow y$ mit $h(z) = g(z)$ für alle $z \in x$ zwar i.A. nicht mehr surjektiv, aber immer noch injektiv. $\square$

Satz 2.22 (Transitivität der Mächtigkeitsrelation). *Seien x, y und z Mengen. Dann folgt aus $x \preceq y$ und $y \preceq z$ die Eigenschaft $x \preceq z$.*

Beweis. Da $x \preceq y$ und $y \preceq z$ gilt, gibt es injektive Funktionen $f: x \rightarrow y$ und $g: y \rightarrow z$. Daher ist die Funktion $g \circ f: x \rightarrow z$ ebenfalls injektiv: Seien $x_1, x_2 \in x$. Angenommen, es gilt $g \circ f(x_1) = g \circ f(x_2)$. Dann ist $g(f(x_1)) = g(f(x_2))$, und wegen der Injektivität von g gilt damit $f(x_1) = f(x_2)$. Aufgrund der Injektivität von f gilt weiters wie erwünscht $x_1 = x_2$. $\square$

Bemerkung. Die Mächtigkeitsrelation $\preceq$ ist reflexiv und nach obigem Satz auch transitiv. Im nächsten Kapitel werden wir beweisen, dass sie darüber hinaus antisymmetrisch ist, es handelt sich also um eine partielle Ordnungsrelation – aber natürlich, wie schon bei der Ordnungsrelation auf den Ordinalzahlen, nur auf einer Meta-Ebene. Denn $\preceq$ ist zwischen allen Mengen definiert; ihr Definitionsbereich wäre also die Menge aller Mengen, die nicht als Element des Universums existiert.

Definition 2.23 (Abzählbarkeit). Eine Menge x heißt *abzählbar*, wenn $x \preceq \omega$ gilt und *abzählbar unendlich*, wenn $x \cong \omega$ gilt. Eine nicht-abzählbare Menge heißt *überabzählbar*.

Bemerkung. Die Mengen $\mathbb{Z}$ und $\mathbb{Q}$ sind abzählbar, während $\mathbb{R}$ und ${}^\omega\omega$ überabzählbar sind. Dies kann z.B. in [4] und [5] nachgelesen werden.

Satz 2.24 (Cantor, 1891). *Sei x eine Menge. Dann gilt $x \prec \mathcal{P}(x)$.*

Beweis. Angenommen, es gäbe eine Bijektion $f: x \rightarrow \mathcal{P}(x)$. Wir betrachten die Teilmenge $z = \{a \in x \mid f(a) \notin a\}$ von x . Falls $f(z) \in z$ gilt, dann erhalten wir – aufgrund der Definition von z – den Widerspruch $f(z) \notin z$. Falls jedoch andersherum $f(z) \notin z$ gilt, haben wir – erneut aufgrund der Definition von z – den Widerspruch $f(z) \in z$. $\square$

Bemerkung. Mithilfe des Satz von Cantor können wir also bereits feststellen, dass es unendlich viele verschiedene Unendlichkeiten gibt, denn die Mengen ω , $\mathcal{P}(\omega)$, $\mathcal{P}(\mathcal{P}(\omega))$, ... haben alle eine andere Mächtigkeit. Eine natürliche Frage ist, welcher Zusammenhang zwischen den überabzählbaren Mengen $\mathbb{R}$, ${}^\omega\omega$ und $\mathcal{P}(\omega)$ besteht: Alle drei sind gleichmächtig, wie z.B. in [5] nachgelesen werden kann.

Die nächste natürliche Frage ist, ob $\mathcal{P}(\omega)$ (bzw. $\mathbb{R}$ oder ${}^\omega\omega$) die „nächste“ Mächtigkeit nach ω ist. Diese Aussage ist, erneut unter der Annahme, dass ZFC widerspruchsfrei ist, unabhängig von ZFC (und damit auch von ZF) und bekannt unter dem Namen *Kontinuumshypothese*. Der Name rührt daher, dass die Menge $\mathbb{R}$ (bzw. $\mathcal{P}(\omega)$ oder ${}^\omega\omega$) im Kontext der Mengenlehre oft als *Kontinuum* bezeichnet wird, weil $\mathbb{R}$ den Zahlenstrahl kontinuierlich abdeckt. Aus dem Englischen kürzen wir die Kontinuumshypothese mit *CH* ab (*continuum hypothesis*) und ihre Verallgemeinerung mit *GCH* (*generalized continuum hypothesis*).

Aussage 2.25 (Kontinuumshypothese, Potenzmengenform). *Es gibt keine Menge z mit der Eigenschaft $\omega \prec z \prec \mathcal{P}(\omega)$.*

Aussage 2.26 (Verallgemeinerte Kontinuumshypothese, Potenzmengenform). *Für jede unendliche Menge x gilt: Es gibt keine Menge z mit der Eigenschaft $x \prec z \prec \mathcal{P}(x)$.*

2.5 Kardinalzahlen

Während Ordinalzahlen uns ermöglicht haben, die (Wohl-)Ordnungsrelation einer Menge in einem gewissen Sinne zu „messen“, wollen wir nun dasselbe mit dem Begriff der Größe einer Menge, also ihrer Mächtigkeit, erreichen. Dafür führen wir den Begriff der *Kardinalzahl* ein.

Definition 2.27 (Kardinalzahl). Eine Ordinalzahl α heißt *Kardinalzahl*, wenn alle kleineren Ordinalzahlen weniger mächtig als α sind, d.h. es gilt $\forall \beta \in \alpha \ \beta \prec \alpha$.

Bemerkung. Die kleinste unendliche Kardinalzahl ist ω . Wenn wir ω im Kontext von Kardinalität betrachten, bezeichnen wir dieselbe Menge in der Regel als $\aleph_0$ (dabei ist $\aleph$, sprich *Aleph*, der erste Buchstabe des hebräischen Alphabets). Diesem Schema folgend heißt die zweitkleinste unendliche Kardinalzahl $\aleph_1$, und so weiter. Sprechen wir von einer allgemeinen Kardinalzahl, so verwenden wir in der Regel den Buchstaben κ .

Definition 2.28 (Nachfolger einer Kardinalzahl). Sei κ eine Kardinalzahl. Dann ist der *Nachfolger* von κ , in Zeichen κ^+ , definiert durch die kleinste als κ mächtigere Kardinalzahl.

Wenn wir in ZFC arbeiten, also das Auswahlaxiom annehmen, so können wir jeder Menge eine eindeutige gleichmächtige Kardinalzahl zuordnen und damit die Mächtigkeit einer Menge für sich alleine betrachtet definieren, ohne auf den Vergleich zweier Mengen angewiesen zu sein.

Definition 2.29 (Mächtigkeit). Sei x eine wohlordenbare Menge. Dann bezeichnen wir die zu x gleichmächtige Kardinalzahl κ als *Mächtigkeit* oder *Kardinalität* von x , in Zeichen $\kappa = |x|$.

Bemerkung. In ZFC existiert $|x|$ für jede Menge x . Es gilt dann $x \prec y \leftrightarrow |x| < |y|$, wobei $<$ wieder die Ordnung der Ordinalzahlen meint. Außerdem haben wir natürlich $|\kappa| = \kappa$ für alle Kardinalzahlen κ .

Definition 2.30 (Mächtigkeit der Potenzmenge). Sei x eine Menge, sodass sich $\mathcal{P}(x)$ wohlordnen lässt. Dann schreiben wir $2^{|x|}$ für $|\mathcal{P}(x)|$.

Mit Kardinalzahlen können wir alternative Versionen der Kontinuumshypothese beschreiben.

Aussage 2.31 (Kontinuumshypothese, Kardinalzahlform). Die Kardinalzahl $2^{\aleph_0}$ existiert und es gilt $\aleph_1 = 2^{\aleph_0}$.

Aussage 2.32 (Verallgemeinerte Kontinuumshypothese, Kardinalzahlform). Für jede unendliche Kardinalzahl κ existiert die Kardinalzahl 2^κ und es gilt $\kappa^+ = 2^\kappa$.

Für den Zweck dieser Arbeit kürzen wir die Kardinalzahlformen mit CH^* bzw. GCH^* ab. Unter der Annahme des Auswahlaxioms sind diese zu den ursprünglichen äquivalent.

Satz 2.33 (Äquivalenz der Varianten der Kontinuumshypothesen). In ZFC gilt: Wenn x eine unendliche Menge ist, dann gibt es genau dann keine Menge z mit der Eigenschaft $x \prec z \prec \mathcal{P}(x)$, wenn $|x|^+ = 2^{|x|}$ gilt.

Beweis. Sei x eine unendliche Menge. Es gilt dann

$$\begin{aligned}
& |x|^+ = 2^{|x|} \\
& \Leftrightarrow |x|^+ \geq 2^{|x|} \wedge |x|^+ \leq 2^{|x|} \\
& \Leftrightarrow |x|^+ \geq 2^{|x|} \quad \text{Satz von Cantor (Satz 2.24)} \\
& \Leftrightarrow \forall \kappa \ (\kappa \text{ ist Kardinalzahl} \rightarrow (\kappa > |x| \rightarrow \kappa \geq 2^{|x|})) \\
& \Leftrightarrow \forall \kappa \ (\kappa \text{ ist Kardinalzahl} \rightarrow (\neg(\kappa > |x|) \vee \kappa \geq 2^{|x|})) \\
& \Leftrightarrow \forall \kappa \ (\kappa \text{ ist Kardinalzahl} \rightarrow (\kappa \leq |x| \vee \kappa \geq 2^{|x|})) \\
& \Leftrightarrow \forall z \ (|z| \leq |x| \vee |z| \geq 2^{|x|}) \quad \text{Auswahlaxiom} \\
& \Leftrightarrow \forall z \ (|z| \leq |x| \vee |z| \geq |\mathcal{P}(x)|) \quad \text{Definition 2.30} \\
& \Leftrightarrow \nexists z \ (|x| < |z| \wedge |z| < |\mathcal{P}(x)|) \\
& \Leftrightarrow \nexists z \ x \prec z \prec \mathcal{P}(x).
\end{aligned}$$

Daraus folgt die Äquivalenz der Varianten der Kontinuumshypothesen. $\square$

Ohne dem Auswahlaxiom kann es Modelle geben, in denen CH wahr ist, aber CH* falsch. Andersherum lässt sich jedoch leicht zeigen, dass die Kardinalzahlform CH* auch ohne Auswahlaxiom die Potenzmengenform CH impliziert. Für die verallgemeinerten Varianten gilt, dass die Potenzmengenform GCH die Kardinalzahlform GCH* impliziert, wie durch den Satz von Sierpiński (Satz 3.4) ersichtlich ist, den wir im nächsten Kapitel beweisen.

2.6 Spiele

Vergleiche für diesen Abschnitt [8]. Wir wollen nun die notwendige Theorie für das Determiniertheitsaxiom, der dritten Aussage des Titels dieser Arbeit, einführen. Zunächst betonen wir dafür die folgende Definition.

Definition 2.34 (*n*-faches kartesisches Produkt). Sei X eine Menge. Dann definieren wir das *n*-fache kartesische Produkt X^n für alle natürlichen n rekursiv über

- $X^0 = \emptyset$,
- $X^1 = X$ und
- $X^{n+1} = X^n \times X$.

Bemerkung. Es lässt sich leicht nachprüfen, dass die Menge X^n gerade aus den *n*-Tupeln besteht, die wir in Definition 2.3 eingeführt haben. Insbesondere entspricht X^n der Menge aller Folgen der Länge *n* mit Elementen aus X . Um die Menge aller endlichen Folgen mit Elementen aus X zu beschreiben, führen wir die Abkürzung $^{<\omega}X = \bigcup_{n \in \omega} X^n$ ein. Zuletzt halten wir noch fest, dass wir die Menge aller unendlichen Folgen mit Elementen aus X – gelegentlich wird diese mit X^ω referenziert – durch die Funktionenmenge $^\omega X$ im Grunde bereits definiert haben, wenn wir die entsprechenden Funktionen als *Folgen* mit Elementen aus X auffassen. Dies erleichtert uns die Notation, da wir z.B. x_i und $(x_0, x_1, \dots)$ statt $x(i)$ und $\{(0, x(0)), (1, x(1)), \dots\}$ schreiben können.

Wir benötigen nun den Begriff eines *Spiele*, wobei wir in dieser Arbeit ausschließlich unendliche 2-Spieler-Spiele betrachten. Dabei wählt im ersten Zug Spieler I ein Element aus einer vorgegeben Menge X und im zweiten Zug Spieler II. Abwechselnd spielen beide Spieler nach demselben Muster unendlich lange weiter, was bedeutet, dass wir jede entstehende Spielrunde als eine unendliche Folge an Elementen von X auffassen können, d.h. als Element von $^\omega X$. Nun müssen wir nur mehr beschreiben, wann ein Spieler das Spiel gewinnt. Dafür wählen wir eine Teilmenge A von $^\omega X$ und sagen für eine Spielrunde $x \in {}^\omega X$, dass Spieler I gewinnt, wenn $x \in A$ ist; andernfalls gewinnt Spieler II.

Definition 2.35 (Spiel). Seien X und A Mengen. Dann heißt $G_X(A)$ *Spiel* auf der Menge X mit der *Gewinnmenge* A , und jedes Element $x \in {}^\omega X$ nennen wir eine *Spielrunde*.

Als nächstes führen wir den Begriff einer *Strategie* ein. Eine Strategie ist eine Funktion, die einem der beiden Spieler ausgehend von den bereits gespielten Zügen den nächsten zu spielenden Zug vorschlägt.

Definition 2.36 (Strategie, Gewinnstrategie). Sei $G_X(A)$ ein Spiel. Dann heißt eine Funktion $\sigma: \bigcup_{n \in \omega} X^{2n} \rightarrow X$ *Strategie für Spieler I*. Sei $y = (y_0, y_1, \dots) \in {}^\omega X$ eine Folge, die die Züge von Spieler II beschreibt. Unter $\sigma * y$ verstehen wir diejenige Spielrunde, bei der Spieler I gemäß der Strategie σ und Spieler II genau die Züge aus y gespielt hat, d.h. $\sigma * y = (\sigma(\emptyset), y_0, \sigma(\sigma(\emptyset), y_0), y_1, \dots)$. Die Strategie σ heißt *Gewinnstrategie*, wenn $\{\sigma * y \mid y \in {}^\omega X\} \subseteq A$.

Analog definieren wir diese Begriffe für Spieler II, d.h. eine *Strategie für Spieler II* ist eine Funktion $\tau: \bigcup_{n \in \omega} X^{2n+1} \rightarrow X$. Sei $z = (z_0, z_1, \dots) \in {}^\omega X$ eine Folge, die die Züge von Spieler I beschreibt. Unter $z * \tau$ verstehen wir diejenige Spielrunde, bei der Spieler II gemäß der Strategie τ und Spieler I genau die Züge aus z gespielt hat, d.h. $z * \tau = (z_0, \tau(z_0), z_1, \tau(z_0, \tau(z_0), z_1), \dots)$. Die Strategie τ heißt *Gewinnstrategie*, wenn $\{z * \tau \mid z \in {}^\omega X\} \subseteq X \setminus A$.

Definition 2.37 (Determiniertheit). Wir nennen ein Spiel $G_X(A)$ *determiniert*, wenn einer der beiden Spieler eine Gewinnstrategie hat. Im Falle von $X = \omega$ bezeichnen wir eine Menge $A \subseteq {}^\omega \omega$ als *determiniert*, wenn das zugehörige Spiel $G_\omega(A)$ determiniert ist.

Beispiel. Gale und Stewart haben bewiesen, dass (bezüglich der weiter unten definierten Topologie) offene und abgeschlossene Teilmengen von ${}^\omega \omega$ determiniert sind.

Aussage 2.38 (Determiniertheitsaxiom). *Jede Teilmenge von ${}^\omega \omega$ ist determiniert.*

Bemerkung. Wir kürzen das Determiniertheitsaxiom mit *AD* (von *axiom of determinacy* aus dem Englischen) ab. Im Gegensatz zum Auswahlaxiom und den Varianten der Kontinuumshypothese wissen wir nicht, ob das Determiniertheitsaxiom widerlegbar oder unabhängig von ZF ist – nur eines ist es, vorausgesetzt ZF ist widerspruchsfrei, sicher nicht: beweisbar. Dies werden wir durch einen Satz von Gale und Stewart (Satz 3.9) beweisen.

2.7 Topologie

Vergleiche für diesen Abschnitt [8] und [9]. Um im nächsten Kapitel einige Aussagen über Spiele beweisen zu können, benötigen wir Begriffe aus der Topologie. Wir geben zunächst eine kurze, allgemeine Einführung in topologische Begriffe und beschäftigen uns dann mit dem Spezialfall einer Topologie auf dem Raum ${}^\omega X$ für eine Menge X .

Definition 2.39 (Topologie). Sei eine Menge X gegeben. Eine Menge $\mathcal{T} \subseteq \mathcal{P}(X)$ heißt *Topologie* für X , wenn folgende Eigenschaften erfüllt sind:

- $\emptyset \in \mathcal{T}$ und $X \in \mathcal{T}$,
- die Vereinigung beliebig vieler Mengen aus $\mathcal{T}$ ist in $\mathcal{T}$ enthalten, und
- der Schnitt endlich vieler Mengen aus $\mathcal{T}$ ist in $\mathcal{T}$ enthalten.

Wir bezeichnen dann die Elemente $x \in \mathcal{T}$ als *offene* Mengen, und eine Menge $y \subseteq X$ heißt *abgeschlossen*, wenn ihr Komplement $X \setminus y$ offen ist.

Definition 2.40 (Isoliertheit). Sei eine Menge X mit einer Topologie $\mathcal{T}$ gegeben, und sei $A \subseteq X$. Dann heißt ein Element $x \in A$ *isoliert in A* , wenn es eine offene Menge $y \in \mathcal{T}$ mit $x \in y$ gibt, sodass $A \cap y = \{x\}$.

Nun wenden wir uns dem konkreten Fall ${}^\omega X$ für eine Menge X zu.

Definition 2.41 (Anfangsstück und Fortsetzung). Sei eine nichtleere Menge X gegeben und sei $a \in {}^{<\omega} X$ und $b \in {}^\omega X$. Dann bezeichnen wir a als *Anfangsstück* von b bzw. b als *Fortsetzung* von a , in Zeichen $a \subseteq b$, wenn es ein $n \in \omega$ gibt, sodass $(b_i)_{i \leq n} = a$ gilt.

Definition 2.42 (Topologie für ${}^\omega X$). Für $s \in {}^{<\omega} X$ definieren wir $\mathcal{O}(s) = \{x \in {}^\omega X \mid s \subset x\}$ als *einfache offene Mengen*. Dann definieren wir durch

$$\mathcal{T} = \{A \in \mathcal{P}(X) \mid \exists S \text{ mit } S \subseteq {}^{<\omega} X \text{ und } A = \bigcup_{s \in S} \mathcal{O}(s)\}$$

eine Topologie auf ${}^\omega X$.

Bemerkung. Unsere Topologie für ${}^\omega X$ ist nicht „zufällig“ entstanden. Es handelt sich genau um die *Produkttopologie* von ${}^\omega X$, wobei wir auf X die *diskrete Topologie* festsetzen, d.h. $\mathcal{T}_X = \mathcal{P}(X)$. Die Produkttopologie eines kartesischen Produktes $\prod X_i$ ist dabei der Schnitt aller Topologien $\mathcal{S}$ für $\prod X_i$, die alle Urbilder von offenen Mengen in der Topologie $\mathcal{T}_{X_j}$ unter den Projektionen $p_j: \prod X_i \rightarrow X_j$ enthalten. Wir fassen dafür also ${}^\omega X$ als unendliches kartesisches Produkt mit den Projektionen $p_j: {}^\omega X \rightarrow X$ mit $p_j(x) = x(j)$ auf. Da wir dabei auf X die diskrete Topologie festsetzen, sind alle Teilmengen von X offene Mengen, d.h. die Urbilder aller Teilmengen von X müssen unter jedem p_j in der Produkttopologie von ${}^\omega X$ enthalten sein. Das bedeutet insbesondere, dass die Menge aller Folgen $(x_i)_i$, die an der Stelle j einen bestimmten Wert annehmen, offen sein muss. Durch die Vereinigungseigenschaft von Topologien ist damit also für jedes $s \in {}^{<\omega} X$ die Menge $\mathcal{O}(s)$ offen. Alle nach obiger Definition offenen Mengen von ${}^\omega X$ sind also auch offene Mengen in der Produkttopologie von ${}^\omega X$. Die andere Richtung lässt sich ebenso beweisen, weshalb es sich bei unserer definierten Topologie tatsächlich um die Produkttopologie handelt (und damit insbesondere um eine wohldefinierte Topologie).

3 Resultate der Mengenlehre

Alle Beweise führen wir, sofern nicht anders angegeben, in ZF. Vergleiche für die ersten drei Abschnitte [4].

3.1 Das Cantor-Bernstein-Schröder-Theorem

Wir beginnen dieses Kapitel mit einem Beweis für das Theorem von Cantor-Bernstein-Schröder (oft abgekürzt als CBS-Theorem), das die Antisymmetrie der Mächtigkeitsrelation behauptet.

Satz 3.1 (Cantor-Bernstein-Schröder).

Seien a und b zwei Mengen, die $a \preceq b$ und $b \preceq a$ erfüllen. Dann folgt $a \cong b$.

Beweis. Da $a \preceq b$ gilt, existiert eine injektive Funktion f von a nach b und wegen $b \preceq a$ eine injektive Funktion g von b nach a . Wir bezeichnen ein Element x für $n \in \omega$ als *von a aus in n Schritten erreichbar*, falls es ein Tupel $(x_0, x_1, \dots, x_n)$ derart gibt, dass $x_0 \in a$, $x_{i+1} = f(x_i)$ für alle geraden Indizes $i \geq 0$ und $x_{i+1} = g(x_i)$ für alle ungeraden Indizes $i > 0$ gilt, sowie $x_n = x$ erfüllt ist. Verständlicher ausgedrückt ist ein Element x dann von a aus in n Schritten erreichbar, wenn wir ausgehend von einem Element x_0 aus a mittels der Funktionen f und g genau n -mal zwischen den Mengen a und b hin- und herspringen, bis wir am Ende beim Element x ankommen.

Weiters bezeichnen wir ein Element x als *a -terminierend*, falls es eine natürliche Zahl n gibt, sodass x in n Schritten von a aus erreichbar ist, aber nicht in m Schritten aus a oder b für alle $m > n$. Analog definieren wir *b -terminierend*. Zuletzt heißt ein Element *nicht-terminierend*, falls es weder a - noch b -terminierend ist, d.h. es keine solche maximale Zahl n gibt.

Nun definieren wir die Mengen a_1 , a_2 und a_3 sowie b_1 , b_2 und b_3 wie folgt:

- $a_1 = \{x \in a \mid x \text{ ist } a\text{-terminierend}\}$
- $a_2 = \{x \in a \mid x \text{ ist } b\text{-terminierend}\}$
- $a_3 = \{x \in a \mid x \text{ ist nicht-terminierend}\}$
- $b_1 = \{x \in b \mid x \text{ ist } a\text{-terminierend}\}$
- $b_2 = \{x \in b \mid x \text{ ist } b\text{-terminierend}\}$
- $b_3 = \{x \in b \mid x \text{ ist nicht-terminierend}\}$

Es ist klar, dass jedes Element von a bzw. b zu genau einer der drei Gruppen a -terminierend, b -terminierend und nicht-terminierend zählt: Jedes Element zählt zu mindestens einer der drei Gruppen aufgrund der Definition von „nicht-terminierend“. Andersherum kann kein Element zu mehr als einer der drei Gruppen zählen, denn wäre ein Element gleichzeitig a - und b -terminierend, müsste die maximale natürliche Zahl n , für die das Element terminiert, gleichzeitig gerade und ungerade sein. Daher ist a die disjunkte Vereinigung von a_1 , a_2 und a_3 , sowie b die disjunkte Vereinigung von b_1 , b_2 und b_3 .

Weiters gilt: f bildet a_1 bijektiv auf b_1 ab, g bildet b_2 bijektiv auf a_2 ab, und sowohl f als auch g bilden a_3 und b_3 bijektiv aufeinander ab. Dafür betrachten wir ein beliebiges Element x aus a_1 . Dann ist x für eine maximale natürliche Zahl n in n Schritten aus a erreichbar. Also ist $f(x)$ in $n + 1$ Schritten aus a erreichbar, und für kein größeres m (sonst wäre n nicht maximal für x). Daher ist $f(x)$ ein Element von b_1 . Damit ist die Wohldefiniertheit gezeigt, und da f injektiv ist, fehlt nur mehr die Surjektivität: Dafür betrachten wir ein beliebiges Element y aus b_1 . Dann ist y nach Definition von b_1 in mindestens einem Schritt aus a erreichbar, es gibt also ein $x \in a$ mit $f(x) = y$. Dieses x muss in a_1 liegen, weil y a -terminierend ist.

Analog kann gezeigt werden, dass $g: b_2 \rightarrow a_2$ eine wohldefinierte Bijektion ist. Die Wohldefiniertheit der letzten Funktion folgt aus der Tatsache, dass das Bild eines nicht-terminierenden Elements ebenfalls nicht-terminierend sein muss (sonst würde auch das ursprüngliche Element terminieren). Ebenso bildet f die Menge a_3 surjektiv auf b_3 ab, weil jedes Element y von b_3 in mindestens einem Schritt von a aus erreichbar sein muss (sonst würde es mit dem Wert 0 terminieren). Allerdings bilden a_1 und a_2 ja, wie wir bereits bewiesen haben, nicht nach b_3 ab; daher muss das Urbild von y unter f in a_3 liegen.

Damit ist die Funktion $h: a \rightarrow b$ definiert durch

$$h(x) = \begin{cases} f(x) & \text{wenn } x \in a_1 \text{ oder } x \in a_3 \\ g^{-1}(x) & \text{wenn } x \in a_2 \end{cases}$$

eine Bijektion von a nach b , also gilt $a \cong b$. □

3.2 Der Satz von Hartogs

Wir wenden uns nun dem sogenannten Satz von Hartogs zu. Dieses machte eine Aussage über „beliebig große“ Ordinalzahlen. Zunächst aber folgende Definition.

Definition 3.2. Wir definieren $\mathcal{P}^n(x)$ für jedes natürliche n rekursiv über

- $\mathcal{P}^0(x) = x$,
- $\mathcal{P}^1(x) = \mathcal{P}(x)$ und
- $\mathcal{P}^{n+1}(x) = \mathcal{P}(\mathcal{P}^n(x))$.

Satz 3.3 (Hartogs).

Für jede Menge x gibt es eine Ordinalzahl α , die nicht injektiv auf x abgebildet werden kann. Darüber hinaus kann α so gewählt werden, dass $\alpha \preceq \mathcal{P}^4(x)$ gilt.

Beweis. Wir bezeichnen mit b die Menge aller möglichen Wohlordnungen für Teilmengen von x . Dann ist b eine Menge, weil jedes Element von b eine Wohlordnung ist, d.h. eine binäre Relation, also eine Teilmenge von $x \times x$. Daher ist $b \subseteq \mathcal{P}(x \times x)$, also eine Menge nach dem Aussonderungsaxiomenschema.

Bezeichne F die Funktion, die jeder Wohlordnung y aus b ihre ordnungsisomorphe Ordinalzahl zuordnet. Dann ist F wohldefiniert nach Satz 2.15. Es existiert also eine

Formel ϕ , die wir für das Ersetzungsaxiomenschema verwenden können und damit erhalten wir, dass $\text{im}(F)$ eine Menge ist. Da $\text{im}(F)$ eine Menge von Ordinalzahlen ist, gibt es eine Ordinalzahl α , die größer als alle Ordinalzahlen in $\text{im}(F)$ ist: Da $\bigcup \text{im}(F)$ gleich groß oder größer als alle Elemente von $\text{im}(F)$ ist und nach Satz 2.15 selbst eine Ordinalzahl ist, ist der Nachfolger von $\bigcup \text{im}(F)$ eine Ordinalzahl α mit der gewünschten Eigenschaft.

Die Ordinalzahl α ist dann nicht ordnungsisomorph zu irgendeinem Element von $\text{im}(F)$, kann also insbesondere nicht bijektiv auf irgendeine Teilmenge von x abgebildet werden und damit nicht injektiv auf x selbst.

Für den zweiten Teil definieren wir zunächst β als die kleinste Ordinalzahl mit der Eigenschaft $\beta \not\leq x$ (wie wir bereits bewiesen haben, gibt es mindestens ein solches β). Wir definieren die Funktion

$$f: \beta \rightarrow \mathcal{P}^4(x) \\ \gamma \mapsto \{a \in \mathcal{P}(x \times x) \mid a \text{ ist eine zu } \gamma \text{ ordnungsisomorphe Wohlordung}\}.$$

Dann ordnet f jedem Element von β , also jeder kleineren Ordinalzahl γ , die Menge aller zu γ ordnungsisomorphen Wohlordnungen von Teilmengen von x zu. Die Funktion f ist injektiv, weil unterschiedliche Ordinalzahlen nach Satz 2.15 nicht zueinander ordnungsisomorph sind. Sobald wir gezeigt haben, dass f wohldefiniert ist (also tatsächlich nach $\mathcal{P}^4(x)$ abbildet), folgt daher $\beta \preceq \mathcal{P}^4(x)$.

Wohldefiniertheit:

$$\begin{aligned} f(\gamma) &= \{a \in \mathcal{P}(x \times x) \mid a \text{ ist eine zu } \gamma \text{ ordnungsisomorphe Wohlordung}\} \\ f(\gamma) &\subseteq \mathcal{P}(x \times x) \\ f(\gamma) &\subseteq \mathcal{P}(\mathcal{P}^2(x)) && \text{denn } x \times x \subseteq \mathcal{P}^2(x) \\ f(\gamma) &\in \mathcal{P}(\mathcal{P}(\mathcal{P}^2(x))) \\ f(\gamma) &\in \mathcal{P}^4(x) \end{aligned}$$

Damit ist der Beweis erbracht. □

3.3 Der Satz von Sierpiński

Wir wollen nun das Theorem von Sierpiński beweisen, das wie bereits erwähnt einen Zusammenhang zwischen der verallgemeinerten Kontinuumshypothese und dem Auswahlaxiom herstellt.

Satz 3.4 (Sierpiński, 1947).

In ZF impliziert die verallgemeinerte Kontinuumshypothese das Auswahlaxiom.

Bemerkung. Der Satz von Sierpinski hat einen Einfluss auf Unabhängigkeitsbeweise der verallgemeinerten Kontinuumshypothese und des Auswahlaxioms: Hat man bereits bewiesen, dass ZF das Auswahlaxiom nicht beweisen kann (in Zeichen: $\text{ZF} \not\vdash \text{AC}$), so folgt mithilfe des Satzes ebenso $\text{ZF} \not\vdash \text{GCH}$. Andersherum, aufgrund der Kontraposition, erhalten wir aus $\text{ZF} \not\vdash \neg \text{GCH}$ die Folgerung $\text{ZF} \not\vdash \neg \text{AC}$.

Bevor wir uns dem Beweis zuwenden, führen wir eine Definition und drei Lemmata ein.

Definition 3.5.

Für zwei Mengen x und y definieren wir die Summe als $x + y := (x \times \{0\}) \cup (y \times \{1\})$.

Bemerkung. Es gilt $x + x = (x \times \{0\}) \cup (x \times \{1\}) = x \times \{0, 1\} = x \times 2$.

Lemma 3.6. Seien x, y, z und v Mengen. Dann gilt:

- (a) Aus $x \preceq y$ und $z \preceq v$ folgt $x + z \preceq y + v$.
- (b) Wenn x und y disjunkt sind, gilt $\mathcal{P}(x \cup y) \cong \mathcal{P}(x) \times \mathcal{P}(y)$. Insbesondere gilt $\mathcal{P}(x + y) \cong \mathcal{P}(x) \times \mathcal{P}(y)$ für alle Mengen x und y .

Beweis.

- (a) Nach Voraussetzung gilt $x \preceq y$ und $z \preceq v$, es gibt also injektive Funktionen $f: x \rightarrow y$ und $g: z \rightarrow v$. Wir definieren eine Funktion $h: x + z \rightarrow y + v$ durch

$$h(c) = \begin{cases} (f(c_1), 0) & \text{wenn } c = (c_1, 0) \\ (g(c_1), 1) & \text{wenn } c = (c_1, 1). \end{cases}$$

Die Funktion h ist damit wohldefiniert und darüber hinaus injektiv: Seien dafür $a, b \in x + z$, sodass $h(a) = h(b)$ gilt. Dann ist entweder $h(a) = h(b) = (y_0, 0)$ für ein $y_0 \in y$ oder $h(a) = h(b) = (v_0, 1)$ für ein $v_0 \in v$.

Im ersten Fall folgt, dass $a = (a_1, 0)$ und $b = (b_1, 0)$ mit $a_1, b_1 \in x$ und $y_0 = f(a_1) = f(b_1)$ gilt. Aus der Injektivität von f folgt dann $a_1 = b_1$ und damit $a = b$. Der zweite Fall liefert analog $a = b$.

Damit ist die Injektivität bewiesen, es folgt also $x + z \preceq y + v$.

- (b) Seien x und y zwei disjunkte Mengen.

Wir definieren die Funktion

$$\begin{aligned} f: \mathcal{P}(x) \times \mathcal{P}(y) &\rightarrow \mathcal{P}(x \cup y) \\ (a, b) &\mapsto a \cup b. \end{aligned}$$

Die Funktion f ist wohldefiniert und wir zeigen, dass f darüber hinaus surjektiv und injektiv ist. Damit folgt die gewünschte Aussage.

Surjektivität: Sei $z \in \mathcal{P}(x \cup y)$ beliebig. Dann ist $z \subseteq x \cup y$, d.h. $z = z_1 \cup z_2$ mit $z_1 \subseteq x$ und $z_2 \subseteq y$. Damit ist $(z_1, z_2) \in \mathcal{P}(x) \times \mathcal{P}(y)$ und $f(z_1, z_2) = z$.

Injektivität: Angenommen, es ist $a \cup b = \tilde{a} \cup \tilde{b}$ für $a, \tilde{a} \subseteq x$ und $b, \tilde{b} \subseteq y$. Da x, y disjunkt zueinander sind, ist a zu b , a zu $\tilde{b}$, $\tilde{a}$ zu b sowie $\tilde{a}$ zu $\tilde{b}$ disjunkt. Entweder a ist leer, woraus gleich $a \subseteq \tilde{a}$ folgt, oder a ist nichtleer. Sei dann $c \in a$ beliebig. Dann ist insbesondere $c \in a \cup b$, also $c \in \tilde{a} \cup \tilde{b}$. Da jedoch a zu $\tilde{b}$ disjunkt ist, muss $c \in \tilde{a}$ gelten. Es folgt also erneut $a \subseteq \tilde{a}$. Analog gilt $\tilde{a} \subseteq a$ und damit $a = \tilde{a}$. Auf gleiche Weise folgt $b = \tilde{b}$.

Nun zum zweiten Teil: Für alle Mengen x und y folgt, dass $\mathcal{P}(x + y) = \mathcal{P}((x \times \{0\}) \cup (y \times \{1\})) \cong \mathcal{P}(x \times \{0\}) \times \mathcal{P}(y \times \{1\}) \cong \mathcal{P}(x) \times \mathcal{P}(y)$, weil $x \times \{0\}$ und $y \times \{1\}$ disjunkt sind.

□

Lemma 3.7. *Sei x eine Menge. Dann gilt:*

- (a) *Aus $\omega \preceq x$ folgt $x \cong x + 1$ und aus $x \cong x + 1$ folgt weiter $\mathcal{P}(x) \cong \mathcal{P}(x) + \mathcal{P}(x)$.*
- (b) *Wenn es eine Menge z gibt, sodass $x = \mathcal{P}(z \cup \omega)$ ist, dann gilt für jede natürliche Zahl n , dass $\mathcal{P}^n(x) \cong \mathcal{P}^n(x) + \mathcal{P}^n(x)$ ist.*

Beweis.

- (a) Wir beginnen mit dem ersten Teil, dass also unter der Voraussetzung $\omega \preceq x$ die Eigenschaft $x \cong x + 1$ folge.

Da $\omega \preceq x$ erfüllt ist, gibt es nach Satz 2.21 eine bijektive Funktion $h: \omega \rightarrow a$ mit $a \subseteq x$. Es gilt also $\omega \cong a$. Sei nun c eine Menge mit $c \notin x$. Wir definieren nun

$$\begin{aligned} \tilde{h}: \omega &\rightarrow a \cup \{c\} \\ n &\mapsto \begin{cases} c & \text{wenn } n = 0 \\ h(n - 1) & \text{wenn } n > 0. \end{cases} \end{aligned}$$

Dann ist $\tilde{h}$ eine Bijektion von ω nach $a \cup \{c\}$, es gilt also auch $\omega \cong a \cup \{c\}$. Daher folgt $a \cong a \cup \{c\}$.

Damit erhalten wir

$$\begin{aligned} x &= a \cup (x \setminus a) \\ &\cong (a \cup \{c\}) \cup (x \setminus a) && \text{Folgt aus obiger Bijektion} \\ &= x \cup \{c\} \\ &\cong (x \times \{0\}) \cup (\{c\} \times \{1\}) && \text{Einfache Bijektion und } c \notin x \\ &= x + 1. \end{aligned}$$

Nun zum zweiten Teil unter der Annahme, dass $x \cong x + 1$ gelte.

$$\begin{aligned} \mathcal{P}(x) &\cong \mathcal{P}(x + 1) \\ &\cong \mathcal{P}(x) \times \mathcal{P}(1) && \text{Lemma 3.6} \\ &= \mathcal{P}(x) \times \mathcal{P}(\{\emptyset\}) \\ &= \mathcal{P}(x) \times \{\emptyset, \{\emptyset\}\} \\ &= (\mathcal{P}(x) \times \{\emptyset\}) \cup (\mathcal{P}(x) \times \{\{\emptyset\}\}) \\ &= (\mathcal{P}(x) \times \{0\}) \cup (\mathcal{P}(x) \times \{1\}) \\ &= \mathcal{P}(x) + \mathcal{P}(x). && \text{Definition 3.5} \end{aligned}$$

- (b) Für jede Menge a gilt $a \preceq \mathcal{P}(a)$ (betrachte dafür die einfache Injektion $v \mapsto \{v\}$). Rekursiv folgt daraus durch die Transitivität von „ $\preceq$ “ (Satz 2.22) $a \preceq \mathcal{P}^n(a)$ für $n > 0$ und ebenso $a \preceq a = \mathcal{P}^0(a)$ für $n = 0$. Da $\omega \preceq z \cup \omega$ gilt (betrachte als Injektion die Identität), haben wir $\omega \preceq \mathcal{P}^n(z \cup \omega)$.

Daher dürfen wir den eben bewiesenen Teil (a) dieses Lemmas verwenden und erhalten für alle $n \in \omega$: $\mathcal{P}(\mathcal{P}^n(z \cup \omega)) = \mathcal{P}(\mathcal{P}^n(z \cup \omega)) + \mathcal{P}(\mathcal{P}^n(z \cup \omega))$. Nun ist aber nach Voraussetzung $x = \mathcal{P}(z \cup \omega)$, daher haben wir

$$\begin{aligned} \mathcal{P}^n(x) &= \mathcal{P}^n(\mathcal{P}(z \cup \omega)) \\ &= \mathcal{P}(\mathcal{P}^n(z \cup \omega)) \\ &\cong \mathcal{P}(\mathcal{P}^n(z \cup \omega)) + \mathcal{P}(\mathcal{P}^n(z \cup \omega)) \\ &= \mathcal{P}^n(\mathcal{P}(z \cup \omega)) + \mathcal{P}^n(\mathcal{P}(z \cup \omega)) \\ &= \mathcal{P}^n(x) + \mathcal{P}^n(x). \end{aligned}$$

Damit ist der Beweis erbracht. □

Lemma 3.8. *Seien x, y und z Mengen. Dann gilt:*

- (a) *Aus $x \cup y \cong z \times z$ folgt: Es gibt eine surjektive Funktion $h: x \rightarrow z$, oder es gilt $z \preceq y$.*
(b) *Aus $x + x \cong x$ und $x + y \cong \mathcal{P}(x)$ folgt $\mathcal{P}(x) \preceq y$.*

Beweis.

- (a) Angenommen, es gibt keine Funktion $h: x \rightarrow z$ mit $\text{im}(h) = z$. Da $x \cup y \cong z \times z$ gilt, gibt es eine Bijektion $f: x \cup y \rightarrow z \times z$. Wir betrachten $f[x]$ und bemerken, dass jedes $c \in f[x]$ die Form (c_1, c_2) für $c_1, c_2 \in z$ hat. Definieren wir die Funktion h als Projektion von $f[x]$ auf die erste Koordinate, d.h. $h(c) = c_1$, dann ist $h \circ f$ eine Funktion von $x \cup y$ nach z . Also ist die Einschränkung $(h \circ f)|_x$ eine Funktion von x nach z , die daher nach unserer Annahme nicht surjektiv sein kann, also mindestens ein Element a auslässt. Daher ist das Urbild der Funktion f von (a, c_2) nicht in x , sondern in y enthalten – und zwar für alle c_2 . Das bedeutet, die Urbildmenge $f^{-1}[\{(a, c_2) \mid c_2 \in z\}]$ ist eine Teilmenge von y . Es ist aber $\{(a, c_2) \mid c_2 \in z\} \cong z$, daher gilt $z \preceq y$.

- (b) Nach der ersten Voraussetzung gilt

$$\begin{aligned} x + x &\cong x \\ \mathcal{P}(x + x) &\cong \mathcal{P}(x) \\ \mathcal{P}(x + x) &\cong x + y && \text{Zweite Voraussetzung} \\ \mathcal{P}(x) \times \mathcal{P}(x) &\cong x + y && \text{Lemma 3.6} \\ \mathcal{P}(x) \times \mathcal{P}(x) &\cong (x \times \{0\}) \cup (y \times \{1\}) \end{aligned}$$

Nun verwenden wir Teil (a) dieses Lemmas: Es gibt eine Funktion $h: x \times \{0\} \rightarrow \mathcal{P}(x)$ mit $\text{im}(h) = \mathcal{P}(x)$, oder $\mathcal{P}(x) \preceq y \times \{1\}$. Beide Aussagen sind äquivalent zur einfacheren Form: Es gibt eine Funktion $h: x \rightarrow \mathcal{P}(x)$ mit $\text{im}(h) = \mathcal{P}(x)$, oder $\mathcal{P}(x) \preceq y$. Analog zum Beweis vom Satz von Cantor (Satz 2.24) – nur mit dem Unterschied, dass h nicht injektiv sein muss –, kann es ein solches h nicht geben. Also folgt $\mathcal{P}(x) \preceq y$ wie erwünscht.

□

Wir beweisen nun den Satz von Sierpiński.

Beweis von Satz 3.4. Nach Voraussetzung gilt die verallgemeinerte Kontinuums-hypothese. Wir führen einen Widerspruchsbeweis durch die Annahme, es gäbe eine Menge, die nicht wohlgeordnet werden kann (nach Satz 2.17 folgt dies aus der Negation des Auswahlaxioms). Sei x diese Menge.

Wir definieren $z = \mathcal{P}(x \cup \omega)$. Da x nicht wohlgeordnet werden kann, gibt es keine Ordinalzahl β mit $x \preceq \beta$, denn wir könnten sonst die Wohlordnung von β durch die vorhandene Injektion auf x übertragen. Da weiters $x \preceq z$ gilt (betrachte dafür die Injektion $v \mapsto \{v\}$) und die Mächtigkeitsrelation nach Satz 2.22 transitiv ist, haben wir für alle Ordinalzahlen β die Eigenschaft

$$z \not\preceq \beta. \quad (1)$$

Nun verwenden wir den Satz von Hartogs (Satz 3.3) und bemerken, dass es eine Ordinalzahl α mit den Eigenschaften

$$\alpha \not\preceq z \quad (2)$$

und $\alpha \preceq \mathcal{P}^4(z)$ gibt. Wir zeigen nun anhand der verallgemeinerten Kontinuums-hypothese, dass wir – je nach auftretendem Fall – entweder einen Widerspruch zu (1) oder zu (2) bekommen.

Wir betonen, dass $\alpha \preceq \mathcal{P}^4(z)$ gilt. Wir wollen nun diese Ungleichung Stück für Stück verschärfen, also zunächst zeigen, dass auch $\alpha \preceq \mathcal{P}^3(z)$ gilt. Aus $\alpha \preceq \mathcal{P}^4(z)$ und $\mathcal{P}^3(z) \preceq \mathcal{P}^4(z)$ (Cantor, Satz 2.24) folgt durch Lemma 3.6 (a) $\mathcal{P}^3(z) + \alpha \preceq \mathcal{P}^4(z) + \mathcal{P}^4(z)$. Da $z = \mathcal{P}(x \cup \omega)$ ist, gilt nach Lemma 3.7 (b) $\mathcal{P}^4(z) = \mathcal{P}^4(z) + \mathcal{P}^4(z)$, also folgt $\mathcal{P}^3(z) + \alpha \preceq \mathcal{P}^4(z)$. Nach Lemma 3.6 (a) gilt andererseits $\mathcal{P}^3(z) \preceq \mathcal{P}^3(z) + \alpha$. Insgesamt erhalten wir

$$\mathcal{P}^3(z) \preceq \mathcal{P}^3(z) + \alpha \preceq \mathcal{P}^4(z).$$

Nun verwenden wir die verallgemeinerte Kontinuums-hypothese, welche impliziert, dass entweder $\mathcal{P}^3(z) \cong \mathcal{P}^3(z) + \alpha$ oder $\mathcal{P}^3(z) + \alpha \cong \mathcal{P}^4(z)$ gilt. In ersterem Fall folgt daraus $\alpha \preceq \mathcal{P}^3(z)$, denn es ist $\alpha \preceq \mathcal{P}^3(z) + \alpha$ nach Lemma 3.6. Damit haben wir die Ungleichung erfolgreich verschärft.

Im zweiten Fall folgt $\mathcal{P}^4(z) \preceq \alpha$ vermöge Lemma 3.8 (b). Das impliziert aber durch die Kette $z \preceq \mathcal{P}(z) \preceq \mathcal{P}^2(z) \preceq \mathcal{P}^3(z) \preceq \mathcal{P}^4(z) \preceq \alpha$ und Lemma 3.6 einen Widerspruch zu (1). Daher ist der zweite Fall unmöglich und wir erhalten stets $\alpha \preceq \mathcal{P}^3(z)$. Diese Argumentation können wir nun dreimal wiederholen, um schlussendlich $\alpha \preceq z$ zu erhalten und damit einen Widerspruch zu (2). □

3.4 Ein Satz von Gale und Stewart

Vergleiche für diesen und den nächsten Abschnitt [8]. Wir beschäftigen uns nun mit zwei Konsequenzen des Determiniertheitsaxioms, welches wir bereits im vorigen Kapitel eingeführt haben. Zunächst bemerken wir, dass es die Negation des Auswahlaxioms impliziert, denn dies ist die Kontraposition des folgenden Satzes.

Satz 3.9 (Gale, Stewart, 1953). *In ZFC gilt: Es gibt eine Teilmenge von ${}^\omega\omega$, die nicht determiniert ist.*

Beweis. Alle möglichen Spieler-I-Strategien für ein G_ω -Spiel sind Funktionen von der Menge $\bigcup_{n \in \omega} \omega^{2^n}$ nach ω . Da der Definitionsbereich abzählbar ist, ist die Menge aller Strategien gleichmächtig zu ${}^\omega\omega$. Wir arbeiten in ZFC, deswegen lässt sich jede Menge wohlordnen und es gibt für jede Menge eine gleichmächtige Kardinalzahl. Im Falle von ${}^\omega\omega$ ist dies die Kardinalzahl $2^{\aleph_0}$, denn es gilt ${}^\omega\omega \cong \mathcal{P}(\omega)$. Daher gibt es eine bijektive Funktion von $2^{\aleph_0}$ in die Menge aller Spieler-I-Strategien. Wir können also eine Aufzählung $(\sigma_\alpha)_{\alpha \in 2^{\aleph_0}}$ definieren. Analog definieren wir $(\tau_\alpha)_{\alpha \in 2^{\aleph_0}}$ für eine Aufzählung aller Spieler-II-Strategien.

Wir definieren nun rekursiv zwei Teilmengen $A = \{a_\alpha \mid \alpha \in 2^{\aleph_0}\}$ und $B = \{b_\alpha \mid \alpha \in 2^{\aleph_0}\}$ von ${}^\omega\omega$. Dabei beziehen wir uns auf den Rekursionssatz für Ordinalzahlen und definieren in allen drei Fällen ($\alpha = \emptyset$, α ist Nachfolgerordinalzahl, α ist Limesordinalzahl) a_α und b_α nach demselben Prinzip. Sei also $\alpha \in 2^{\aleph_0}$. Angenommen, für alle $\beta \in \alpha$ haben wir a_β und b_β bereits definiert. Dann wählen wir ein y aus ${}^\omega\omega$ passend, sodass

$$b_\alpha = \sigma_\alpha * y \text{ und } b_\alpha \notin \{a_\beta \mid \beta \in \alpha\}$$

gilt. Dies ist immer möglich, weil $\{a_\beta \mid \beta \in \alpha\} \preceq \alpha \prec 2^{\aleph_0} \cong \{\sigma_\alpha * y \mid y \in {}^\omega\omega\}$ erfüllt ist, wobei wir verwendet haben, dass $2^{\aleph_0}$ eine Kardinalzahl ist und die Funktion $y \mapsto \sigma_\alpha * y$ injektiv ist. Analog wählen wir ein z aus ${}^\omega\omega$, sodass

$$a_\alpha = z * \tau_\alpha \text{ und } a_\alpha \notin \{b_\beta \mid \beta \in \alpha + 1\}$$

gilt. Nach Definition sind A und B disjunkt. Da es nach Konstruktion von B für jede mögliche Spieler-I-Strategie eine Menge y gibt, die Spieler II spielen kann, sodass die entstehende Spielrunde in B (und damit aufgrund der Disjunktheit nicht in A) liegt, hat Spieler I keine Gewinnstrategie. Ebenso hat aber auch Spieler II keine Gewinnstrategie, weil es für jede Spieler-II-Strategie eine Menge z gibt, die Spieler I spielen kann, sodass die entstehende Spielrunde in A liegt. Daher haben beide Spieler keine Gewinnstrategie für das Spiel $G_\omega(A)$, d.h. die Menge A ist nicht determiniert. $\square$

3.5 Der Satz von Davis

Nun zeigen wir, dass das Determiniertheitsaxiom die Kontinuumshypothese in der Potenzmengenform impliziert. Dafür benötigen wir eine Definition und ein Lemma.

Definition 3.10 (Perfektheit). Eine Teilmenge x von ${}^\omega\omega$ heißt *perfekt*, wenn x nichtleer und abgeschlossen ist sowie keine isolierten Punkte besitzt. Wir sagen, dass eine Teilmenge y von ${}^\omega\omega$ die *Perfekte-Mengen-Eigenschaft* hat, wenn y abzählbar ist oder eine perfekte Teilmenge hat.

Bemerkung. Die Eigenschaften der Abgeschlossenheit und der Nichtexistenz isolierter Punkte sind in einem gewissen Sinne komplementär zueinander: Sofern wir über eine Metrik verfügen, ist die Abgeschlossenheit einer Menge x äquivalent dazu, dass jede konvergierende Folge $(y_n)_n$ an Elementen aus x einen Grenzwert y in x hat, und die Nichtexistenz isolierter Punkte ist äquivalent dazu, dass es für jedes y aus x eine Folge $(y_n)_n$ mit Elementen aus x gibt, sodass $(y_n)_n$ gegen y konvergiert. Eine Menge x ist also genau dann perfekt, wenn sie nichtleer ist und die Funktion f , welche jeder konvergierenden Folge an Elementen aus x ihren Grenzwert zuordnet, die Eigenschaft $\text{im}(f) = x$ erfüllt.

Lemma 3.11. *Sei eine x eine perfekte Menge. Dann gilt $x \cong {}^\omega\omega$.*

Beweis. Sei $x \subseteq {}^\omega\omega$ perfekt. Wir zeigen, dass es eine injektive Funktion $f: {}^\omega 2 \rightarrow x$ gibt, wonach ${}^\omega 2 \preceq x$ gilt. Es lässt sich leicht zeigen, dass ${}^\omega 2 \cong {}^\omega\omega$ gilt, daher erhalten wir dann ${}^\omega\omega \preceq x$. Aufgrund von $x \subseteq {}^\omega\omega$ für alle perfekten Mengen gilt aber auch $x \preceq {}^\omega\omega$, daher haben wir ${}^\omega\omega \preceq x \preceq {}^\omega\omega$, also nach Satz 3.1 (CBS) wie erwünscht $x \cong {}^\omega\omega$.

Sei also ein Element $(y_n)_n \in {}^\omega 2$ gegeben. Wir definieren den Funktionswert $f((y_n)_n)$ nun rekursiv über eine Hilfsfolge $(A_n)_n$, von der jeder Eintrag einer endlichen Folge natürlicher Zahlen entspricht. Der Funktionswert $f((y_n)_n)$ wird schlussendlich aus der Konkatenation aller Folgen A_i bestehen.

Angenommen, wir haben alle Folgen A_i für $i < j$ bereits bestimmt. Es sei dann m der kleinste Index derart, dass es zwei Fortsetzungen von $A_0 \frown \dots \frown A_{j-1}$ in x gibt, die sich im m -ten Folgenglied unterscheiden. Es gibt immer einen solchen Index m , weil x keine isolierten Punkte hat und nichtleer ist. Von allen möglichen Fortsetzungen $(d_n)_n$ nehmen wir dabei zunächst die nach lexikographischer Sortierung kleinste und bezeichnen sie mit $(b_n)_n$. Nun nehmen wir von allen möglichen Fortsetzungen $(d_n)_n$ diejenige in lexikographischer Sortierung kleinste, die $d_m \neq b_m$ erfüllt und bezeichnen sie mit $(c_n)_n$.

Sei k die Länge der Folge $A_0 \frown \dots \frown A_{j-1}$. Wir setzen nun

$$A_j = \begin{cases} (b_i)_{k < i \leq m} & \text{wenn } y_j = 0 \\ (c_i)_{k < i \leq m} & \text{wenn } y_j = 1. \end{cases}$$

Aufgrund der Abgeschlossenheit von x ist die Konkatenation aller A_i ein Element von x und wir setzen

$$f((y_n)_n) = A_0 \frown A_1 \frown A_2 \frown \dots$$

Die Funktion f ist damit wohldefiniert. Da es für zwei verschiedene $(y_n)_n$ und $(\tilde{y}_n)_n$ aus ${}^\omega 2$ einen Index m gibt, sodass $y_m \neq \tilde{y}_m$, ist $a_m \neq \tilde{a}_m$ und damit $f((y_n)_n) \neq f((\tilde{y}_n)_n)$, die Funktion f ist also injektiv. $\square$

Wir können nun den Begriff der Determiniertheit mit dem von perfekten Mengen verbinden, wie der folgende Satz zeigt.

Satz 3.12 (Davis). *Aus dem Determiniertheitsaxiom folgt, dass alle Teilmengen von ${}^\omega\omega$ über die Perfekte-Mengen-Eigenschaft verfügen. Insbesondere gilt also die Kontinuums-hypothese in Potenzmengenform (CH).*

Beweis. Es gelte das Determiniertheitsaxiom. Für diesen Beweis betrachten wir eine adaptierte Spielvariante, die wir mit $G_2^*(A)$ für $A \subseteq {}^\omega 2$ bezeichnen. Dabei spielt Spieler I immer eine (möglicherweise leere) endliche Folge an Nullen und Einsen $s_i \in {}^{<\omega}2$ und Spieler II immer eine 0 oder eine 1, d.h. ein $k_i \in 2$. Eine mögliche Spielrunde x sieht also beispielsweise so aus:

$$x = (s_0, k_1, s_2, k_3, s_4, k_5, s_6, k_7, \dots) = ((1, 0, 1), 0, (), 1, (0), 1, (1, 1, 0, 1), 0, \dots).$$

Wir bezeichnen die durch Konkatenation aller Folgenglieder von x (wobei wir die Spielzüge von Spieler II als Folgen der Länge 1 auffassen) entstehende unendliche Folge aus ${}^\omega 2$ mit $\hat{x}$. In unserem Beispiel wäre

$$\hat{x} = s_0 \frown (k_1) \frown s_2 \frown (k_3) \frown s_4 \frown (k_5) \frown s_6 \frown (k_7) \frown \dots = (1, 0, 1, 0, 1, 0, 1, 1, 1, 0, 1, 0, \dots).$$

Liegt $\hat{x}$ in A , so gewinnt Spieler I; ansonsten gewinnt Spieler II.

Die Begriffe der Strategie bzw. Gewinnstrategie definieren wir analog zur Definition bei den „klassischen“ Spielen.

Schritt 1: Wir zeigen zunächst, dass eine Teilmenge $A \subseteq {}^\omega 2$ genau dann eine perfekte Teilmenge hat, wenn Spieler I eine Gewinnstrategie für das Spiel $G_2^*(A)$ hat.

Hinrichtung: Sei $A \subseteq {}^\omega 2$ mit einer perfekten Teilmenge B . Wir konstruieren nun eine Spieler-I-Gewinnstrategie. Sei dafür

$$S = \{(x_i)_{i \leq n} \mid x \in B \wedge n \in \omega\}.$$

Spieler I beginnt mit einem $s_0 \in S$, sodass sowohl $s_0 \frown (0)$ als auch $s_0 \frown (1)$ in S sind. Ein solches s_0 gibt es immer, weil B keine isolierten Punkte hat und nichtleer ist. Nun definieren wir die Gewinnstrategie rekursiv: Angenommen, beide Spieler haben bereits die Züge $p = (s_0, k_1, \dots, s_{2n}, k_{2n+1})$ mit $s_i \in {}^{<\omega}2$ und $k_i \in 2$ gespielt, sodass insbesondere $\hat{p} = s_0 \frown (k_1) \frown \dots \frown s_{2n} \frown (k_{2n+1})$ in S liegt. Spieler I wählt dann ein s_{2n+2} derart, dass sowohl $\hat{p} \frown (0)$ als auch $\hat{p} \frown (1)$ in S liegen. Erneut gibt es ein solches s_{2n+2} immer, weil B keine isolierten Punkte hat. Da B abgeschlossen ist, liegt die Konkatenation $\hat{x}$ zur jeder nach dieser Strategie entstehenden Spielrunde x immer in B , also insbesondere auch in A , was die Strategie zu einer Gewinnstrategie macht.

Rückrichtung: Sei $A \subseteq {}^\omega 2$ und σ eine Spieler-I-Gewinnstrategie für das Spiel $G_2^*(A)$. Die Menge

$$B = \{\sigma * y \mid y \in {}^\omega 2\}$$

ist dann einerseits nach Definition eine Teilmenge von A , weil σ eine Gewinnstrategie ist, und andererseits perfekt:

- B ist nichtleer, weil ${}^\omega 2$ nichtleer ist.
- B ist abgeschlossen, weil das Komplement ${}^\omega 2 \setminus B$ offen ist: Sei dazu $x \in {}^\omega 2 \setminus B$ beliebig. Dann hat Spieler I in der Spielrunde x an irgendeinem Index m nicht gemäß σ gespielt. Alle Fortsetzungen der endlichen Folge $(x_i)_{i \leq m}$ sind also auch Spielrunden, in denen Spieler I nicht gemäß σ gespielt hat, daher sind alle diese Fortsetzungen in ${}^\omega 2 \setminus B$ enthalten. Das bedeutet aber gerade, dass die Menge $\mathcal{O}((x_i)_{i \leq m})$ eine Teilmenge von ${}^\omega 2 \setminus B$ ist. Da offensichtlich $x \in \mathcal{O}((x_i)_{i \leq m})$ gilt und x beliebig war, können wir ${}^\omega 2 \setminus B$ als Vereinigung der Mengen $\mathcal{O}((x_i)_{i \leq m})$ für alle $x \in {}^\omega 2 \setminus B$ schreiben. Also ist ${}^\omega 2 \setminus B$ offen und B damit abgeschlossen.
- B hat keine isolierten Punkte: Angenommen, es gäbe ein isoliertes $x \in B$. Dann gibt es ein $y \in {}^\omega 2$ derart, dass $\sigma * y = x$ gilt. Nach Definition von Isoliertheit gibt es außerdem eine endliche Folge $s = (s_0, k_1, \dots, s_{2n}, k_{2n+1})$ mit $s_i \in {}^{<\omega} 2$ und $k_i \in 2$, sodass $B \cap \mathcal{O}(s) = \{x\}$ gilt. Dann gilt für $\tilde{y} = (\tilde{y}_i)$ mit

$$\tilde{y}_i = \begin{cases} y_i & \text{wenn } i \leq n \\ 1 - y_i & \text{wenn } i > n \end{cases},$$

dass $\sigma * \tilde{y} \in B$ ist und $\sigma * \tilde{y} \in \mathcal{O}(s)$. Also folgt $\sigma * \tilde{y} \in B \cap \mathcal{O}(s)$, und da nach Definition von $\tilde{y}$ die Eigenschaft $\sigma * \tilde{y} \neq x$ gilt, haben wir einen Widerspruch zu $B \cap \mathcal{O}(s) = \{x\}$ gefunden.

Schritt 2: Nun zeigen wir, dass eine Teilmenge $A \subseteq {}^\omega 2$ genau dann abzählbar ist, wenn Spieler II eine Gewinnstrategie für das Spiel $G_2^*(A)$ hat.

Hinrichtung: Sei $A \subseteq {}^\omega 2$ abzählbar und sei $(a_i)_{i \in \omega}$ eine Abzählung von A . Eine Gewinnstrategie für Spieler II liegt dann darin, im $2i+1$ -ten Zug des Spiels ein $k_{2i+1} \in 2$ derart zu spielen, dass $s_0 \frown (k_1) \frown \dots \frown s_{2i} \frown (k_{2i+1})$ kein Anfangsstück von a_i ist.

Rückrichtung: Angenommen, es gibt eine Spieler-II-Gewinnstrategie τ . Sei dann $p = (s_0, k_0, \dots, s_{2i}, k_{2i+1})$ für ein $i \in \omega$ eine angefangene Spielrunde, in der sich Spieler II an τ hält. Erneut bezeichnen wir mit $\hat{p}$ die Konkatenation aller Folgenglieder von p , d.h. $\hat{p} = s_0 \frown (k_1) \frown \dots \frown s_{2i} \frown (k_{2i+1})$. Wir definieren nun

$$D_p = \{x \in {}^\omega 2 \mid \hat{p} \subseteq x \rightarrow \exists s \in {}^{<\omega} 2 \ \hat{p} \frown s \frown (\tau(p \frown s)) \subseteq x\}.$$

Die Menge D_p besteht dann aus all jenen Folgen, für die es, sofern sie mit $\hat{p}$ beginnen, ein $s \in {}^{<\omega} 2$ – quasi als möglicher Zug von Spieler I – derart gibt, dass die Fortsetzung von $\hat{p} \frown s$ anhand der Strategie τ ein gültiger Anfang der Folge bleibt. Es sei nun $P = \{p \in {}^{<\omega} ({}^{<\omega} 2) \mid p \text{ ist eine angefangene Spielrunde, in der sich Spieler II an } \tau \text{ gehalten hat}\}.$

Behauptung: Es gilt $A \subseteq \bigcup_{p \in P} ({}^\omega 2 \setminus D_p)$.

Beweis der Behauptung: Wir betrachten die äquivalente, komplementäre Aussage ${}^\omega 2 \setminus \bigcup_{p \in P} ({}^\omega 2 \setminus D_p) \subseteq {}^\omega 2 \setminus A$. Diese ist weiter äquivalent zu $\bigcap_{p \in P} D_p \subseteq {}^\omega 2 \setminus A$. Sei also $a \in \bigcap_{p \in P} D_p$ beliebig. Wir wollen nun zeigen, dass $a \in {}^\omega 2 \setminus A$ ist. Dafür genügt es, eine Spielrunde x derart zu konstruieren, dass einerseits $\hat{x} = a$ ist und andererseits Spieler II sich in x stets an τ gehalten hat, d.h. es ein $z \in {}^\omega (<{}^\omega 2)$ gibt, sodass $z * \tau = x$ ist. Denn da τ eine Gewinnstrategie ist, muss dann $\hat{x}$ in ${}^\omega 2 \setminus A$ liegen, d.h. es ist wie erwünscht $a = \hat{x} \in {}^\omega 2 \setminus A$.

Sei also x eine Spielrunde, sodass $\hat{x} = a$ ist. Ein solche Spielrunde existiert stets, denn wir könnten z.B. alle Einträge mit geraden Indizes von a als Folgen der Länge 1 und damit Züge von Spieler I auffassen, und die Einträge mit ungerade Indizes als die Züge von Spieler II. Es gilt nun folgendes: Entweder, Spieler II hat in x gemäß τ gespielt – dann sind wir fertig –, oder es gibt einen kleinsten Index $2n + 1$, sodass $x_{2n+1} \neq \tau((x_i)_{i \leq 2n})$ ist. Da jedoch $\hat{x} \in \bigcap_{p \in P} D_p$ ist und der Anfang p der Spielrunde x bis zum Zug $2n - 1$ gemäß τ gespielt wurde, gibt es eine Folge $s \in <{}^\omega 2$ derart, dass $(\hat{p} \frown s \frown \tau(p \frown (s))) \subseteq \hat{x} = a$ erfüllt ist. Wir können also für eine alternative Fortsetzung $\bar{x}$ von p den Spieler I im Zug $2n$ die Folge s spielen lassen und Spieler II im Zug $2n + 1$ gemäß τ . Die restlichen Züge wählen wir so, dass weiterhin $\hat{\bar{x}} = a$ erfüllt ist. Dieses alternative Spiel $\bar{x}$ erfüllt nun die Eigenschaft, bis zum Zug $2n + 1$ gemäß τ gespielt worden zu sein und dennoch $\hat{\bar{x}} = a$ zu erfüllen. Nach demselben Muster können wir nun rekursiv fortfahren, sodass wir eine Spielrunde erhalten, die gemäß τ gespielt wird und deren Konkatenation gleich a ist.

Wir zeigen nun, dass $\bigcup_{p \in P} ({}^\omega 2 \setminus D_p)$ eine abzählbare Menge ist, woraus die erwünschte Abzählbarkeit von A folgt. Da es nur abzählbar viele endliche Spielanfänge p gibt, genügt es zu zeigen, dass alle ${}^\omega 2 \setminus D_p$ abzählbar sind. Tatsächlich sind diese Mengen (aufgrund der Struktur des Spiels auf der Menge 2) sogar einelementig: Sei dafür ein Spielanfang p vorgegeben. Wir bezeichnen die Länge von $\hat{p}$ mit m . Die Menge ${}^\omega 2 \setminus D_p$ besteht dann gerade aus denjenigen Folgen, die mit $\hat{p}$ beginnen und für die es kein $s \in <{}^\omega 2$ gibt, sodass $\hat{p} \frown s \frown (\tau(p \frown (s)))$ ein Anfang der Folge wäre. Jedes Element $x = (x_i)_i$ von ${}^\omega 2 \setminus D_p$ muss also bis zum Index m mit $\hat{p}$ übereinstimmen. Für den $m + 1$ -ten Eintrag (der an der Stelle m erfolgt) gilt dann aber $x_m = 1 - \tau(p \frown ())$, denn für die einzig andere Möglichkeit für x_m gäbe es mit der Wahl $s = ()$ eine Folge, die Spieler I spielen könnte, sodass doch $\hat{p} \frown s \frown (\tau(p \frown (s))) \subseteq x$ wäre. Diese Argumentation können wir für jeden weiteren Eintrag rekursiv wiederholen: Für alle $l > m$ gilt $x_l = 1 - \tau(p \frown (x_m, \dots, x_{l-1}))$. Also gibt es exakt ein Element in ${}^\omega 2 \setminus D_p$.

Aus Schritt 1 und 2 folgt also, dass eine Teilmenge $A \subseteq {}^\omega 2$ genau dann die Perfekte-Mengen-Eigenschaft hat, wenn das Spiel $G_2^*(A)$ determiniert ist.

Schritt 3: Nun gilt es, den Fall $X = 2$ zum Fall $X = \omega$ zu erweitern. Dafür benutzen wir einen Homöomorphismus $\pi: {}^\omega \omega \rightarrow {}^\omega 2 \setminus E$, wobei E die Menge aller Folgen in ${}^\omega 2$ ist, die ab irgendeinem Index konstant sind. Ein Homöomorphismus ist eine stetige,

bijektive Funktion, deren inverse Funktion auch stetig ist. Stetigkeit bedeutet in diesem Fall, dass das Urbild jeder offenen Menge wieder offen ist.

Sei $x \in {}^\omega\omega$. Wir „kodieren“ dabei für jedes $i \in \omega$ die Zahl x_i durch eine Folge an Einsen oder Nullen, die genau $x_i + 1$ lang ist, wobei wir nach jeder Zahl zwischen Einsen oder Nullen wechseln. Die Folge $(0, 1, 2, 3, \dots)$ wird so zur Folge $(0, 1, 1, 0, 0, 0, 1, 1, 1, \dots)$. Formal können wir dies beschreiben durch

$$\pi(x) = a_{x_0} \frown b_{x_1} \frown a_{x_2} \frown b_{x_3} \frown \dots,$$

wobei a_{x_i} eine $x_i + 1$ -lange Folge an Nullen und b_{x_i} eine $x_i + 1$ -lange Folge an Einsen ist. Es lässt sich leicht zeigen, dass π nach dieser Definition ein Homöomorphismus ist und jede Menge $B \subseteq {}^\omega 2$ genau dann über die Perfekte-Mengen-Eigenschaft verfügt, wenn $\pi^{-1}[B] \subseteq {}^\omega\omega$ dies tut.

Sei also $A \subseteq {}^\omega\omega$. Das Determiniertheitsaxiom impliziert dann insbesondere, dass das Spiel $G_2^*(\pi[A])$ determiniert ist, denn wir können jedes G_2^* -Spiel leicht in ein G_ω -Spiel umwandeln – z.B. indem wir jede Folge, die Spieler I spielt, als Binärzahl auffassen und dann mit der zugehörigen natürlichen Zahl identifizieren. Ebenso kann jede Zahl aus 2 , die Spieler II spielt, mit den geraden und ungeraden natürlichen Zahlen identifiziert werden.

Schritt 1 und 2 liefern dann, dass $\pi[A]$ über die Perfekte-Mengen-Eigenschaft verfügt, und damit auch A , womit der Beweis erbracht ist.

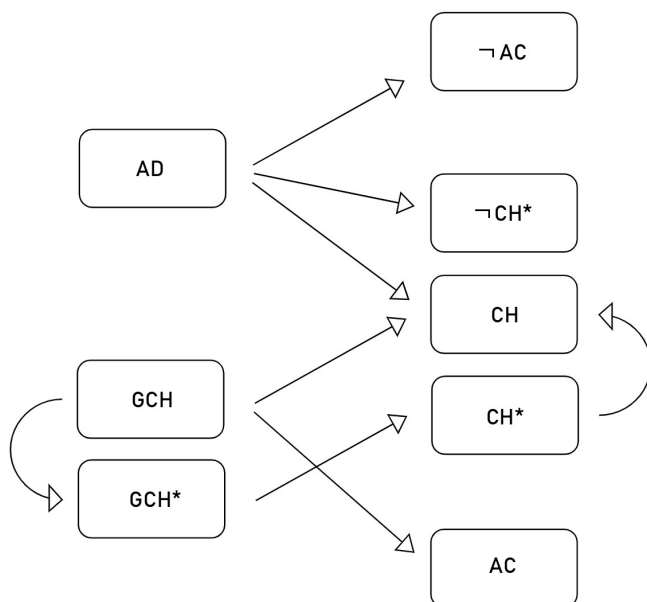
Schritt 4: Zuletzt halten wir fest, dass die Kontinuumshypothese in Potenzmengenform gilt. Angenommen, es gäbe eine Menge z mit $\omega \prec z \prec \mathcal{P}(\omega)$. Da $z \prec \mathcal{P}(\omega)$ und $\mathcal{P}(\omega) \cong {}^\omega\omega$ gilt, gibt es eine Teilmenge x von ${}^\omega\omega$ mit $z \cong x$. Also gilt auch $\omega \prec x \prec \mathcal{P}(\omega)$. Da x jedoch eine Teilmenge von ${}^\omega\omega$ ist, verfügt x nach diesem Beweis über die Perfekte-Mengen-Eigenschaft. Entweder ist x also abzählbar, was ein Widerspruch zu $\omega \prec x$ ist, oder es gibt also eine Teilmenge $y \subseteq x$, die perfekt ist. Dann gilt aber nach Lemma 3.11 $y \cong {}^\omega\omega \cong \mathcal{P}(\omega)$, und damit $\mathcal{P}(\omega) \preceq x$, was ein Widerspruch zu $x \prec \mathcal{P}(\omega)$ ist. $\square$

4 Fazit

Abschließend wollen mittels der folgenden Grafik festhalten, mit welchen Aussagen und Beweisen wir uns in dieser Arbeit beschäftigt haben.

- **AD** Determiniertheitsaxiom *Jede Teilmenge des Kontinuums ist determiniert.*
- **AC** Auswahlaxiom *Für jede Menge nichtleerer paarweiser disjunkter Mengen existiert eine Auswahlmenge.*
- **CH** Kontinuumshypothese in Potenzmengenform *Es gibt keine Menge z , die mächtiger als ω und weniger mächtig als das Kontinuum ist.*
- **CH*** Kontinuumshypothese in Kardinalzahlform *Die Kardinalzahl $2^{\aleph_0}$ existiert und ist die zweitkleinste, unendliche Kardinalzahl.*
- **GCH** Verallgemeinerte Kontinuumshypothese in Potenzmengenform *Für jede unendliche Menge x gibt es keine Menge z mit $x \prec z \prec \mathcal{P}(x)$.*
- **GCH*** Verallgemeinerte Kontinuumshypothese in Kardinalzahlform *Für jede unendliche Kardinalzahl κ existiert die Kardinalzahl 2^κ und es gilt $\kappa^+ = 2^\kappa$.*

Die Pfeile stellen jeweils eine Implikation in ZF dar. Dabei handelt es sich bei $AD \rightarrow \neg AC$ um den Satz von Gale und Stewart (Satz 3.9), bei $AD \rightarrow CH$ um den Satz von Davis (Satz 3.12) und bei $GCH \rightarrow AD$ um den Satz von Sierpiński (Satz 3.4). Letzterer erlaubt uns dank der Äquivalenz der Kontinuumshypothesen unter dem Auswahlaxiom (Satz 2.33) die Folgerung $GCH \rightarrow GCH^*$. Die Implikationen $GCH \rightarrow CH$ sowie $GCH^* \rightarrow CH^*$ sind trivial. Nicht bewiesen haben wir $AD \rightarrow \neg CH^*$ und $CH^* \rightarrow CH$, wobei letzteres eine leichte Übung ist und ersteres als *Satz von Mycielski* bekannt ist und z.B. in [7] nachgelesen werden kann.



Literatur

- [1] Hoffmann, Dirk W. *Grenzen der Mathematik*. Springer Spektrum, 3. Auflage 2018
- [2] Rautenberg, Wolfgang. *Einführung in die mathematische Logik*. Vieweg+Teubner, 3. Auflage 2008.
- [3] Müller, Sandra. *Skript zur Vorlesung Grundzüge der mathematischen Logik*. 2020.
- [4] Smullyan, Raymond M. und Fitting, Melvin. *Set theory and the continuum problem*. Dover Publications, 1. Auflage 2010.
- [5] Deiser, Oliver. *Einführung in die Mengenlehre*. Ohne Verlag, München, 2016.
- [6] Ebbinghaus, Hans-Dieter. *Einführung in die Mengenlehre*. Springer Spektrum, 5. Auflage 2021.
- [7] Kanamori, Akihiro. *The Higher Infinite. Large Cardinals in Set Theory from Their Beginnings*. Springer-Verlag Berlin Heidelberg, 2. Auflage 2009.
- [8] Müller, Sandra. *Determinacy: Lecture Notes For Advanced Topics in Mathematical Logic*. 2020.
- [9] Laures, Gerd und Szymik, Markus. *Grundkurs Topologie*. Springer Spektrum, 3. Auflage 2023.