

Das Antinomien – Problem

Andreas Otte

1987 - 2002

Zusammenfassung

In dem folgenden Text wird wiederum versucht, die Thematik der klassischen Antinomien zu beleuchten. Die Ergebnisse erlauben es außerdem, Cantors 2. Diagonalverfahren sowie Kurt Gödels berühmten Unvollständigkeitssatz in ähnlicher Weise zu behandeln, und weichen zum Teil erheblich von der heute üblichen Schulmeinung ab.

Im einzelnen:

Die Antinomien erweisen sich als Antilogien, als immer falsche Ausdrücke über deren logische Konsequenzen man sich nicht wundern muß. Die Entstehung der Antinomien beruht auf einer naiven Einstellung zur Gültigkeit von Definitionen.

Bei den Diagonalverfahren ergibt sich, daß der verwendete Schluß nach dem *modus tollens* mindestens noch eine weitere ganz wichtige, bisher aber anscheinend unbeachtete Prämisse (die Diagonaldefinition) enthält, deren Widerspruchsfreiheit nicht nachzuweisen ist, da es sich um zirkulär veränderliche Definitionen handelt. Der *modus tollens* weist also nicht eindeutig auf den Verursacher des Widerspruches hin, womit die Beweise durch Diagonalverfahren unschlüssig werden.

Beschränkt man sich auf den Einführungsteil, so bedient sich Gödel bei seinem Unvollständigkeitssatz eines Tricks, der zunächst bewirkt, daß sein Satz nicht als Antilogie anzusehen ist. Unter der Hinzunahme einer weiteren, doch recht zwingenden Voraussetzung, erweist sich der Satz dann aber doch als Antilogie, so daß sich letztlich ergibt, daß der Satz in einem vollständigen, widerspruchsfreien System eine Antilogie, in einem unvollständigen, widerspruchsfreien System dagegen eine Trivialität ist.

Die Antinomienproblematik in Kombination mit dem 2. Cantorschen Diagonalverfahren sowie Gödels Satz ist wieder aktuell geworden, nachdem vor einiger Zeit das Buch „Gödel, Escher, Bach“ von Douglas R. Hofstadter [14] erschienen ist, und sich inzwischen zum Kultbuch entwickelt hat.

1 Einleitung

Unter Antinomien versteht man bekanntlich widerspruchsfreie Ausdrücke eines Kalküls, aus denen sich auf korrekten Wegen (nach den Regeln des jeweiligen Kalküls) Widersprüche herleiten lassen. Zu diesem Thema ist schon viel geschrieben worden, und das mit Recht, denn dieses Problem ist von zentraler Bedeutung, sowohl für die Logik, als auch für die Mathematik. Gäbe es z.B. in der klassischen Logik Antinomien, so würde das den Zusammenbruch der Logik und damit auch der

Mathematik, so wie wir sie bisher kennen, bedeuten. Der Reigen der Antinomiekandidaten und verwandter Probleme in der klassischen Logik reicht von dem so bekannten „Lügner des Epimenides“ über Russells „Menge der Dinge, die sich nicht selbst als Element enthalten“, bis zu Kurt Gödels Aufsatz über „unentscheidbare Sätze der Principia Mathematica und verwandter Systeme“, um nur einige wenige Beispiele zu nennen.

Eine Untersuchung der Antinomien der klassischen Logik ist schon oft vorgenommen worden, aber eine allgemein anerkannte Lösung gibt es bisher, aus welchen Gründen auch immer, nicht. Der Schwerpunkt der Antinomien-Diskussion hat sich daher bald von der Frage ihrer Entstehung hin zur Vermeidung der Antinomien verschoben, also zu einer symptomatischen Vorgehensweise. Hier sind viele interessante Systementwürfe entstanden, insbesondere zahlreiche konstruktivistische Systeme. Eine rein symptomatische Behandlung der Antinomien würde genügen, wenn es sich um Widersprüche in einem beliebigen nicht weiter interessanten Kalkül mit beliebiger Axiomatik handeln würde. Dann würde man eben in einem widersprüchlichen Kalkül arbeiten.

Das Problem ist, daß die Antinomien scheinbar in einem besonderen Kalkül, nämlich der „klassischen Logik“, nachweisbar sind, der sich als adäquate Formalisierung des inhaltlichen Schließens, so wie es bislang anerkannt ist, versteht, und da sollte es keine Widersprüche geben.

Was ist das sogenannte inhaltliche Schließen? Darunter versteht man im allgemeinen das leider oft etwas unscharfe, noch dazu subjektabhängige Gefühl, daß ein Folgerungszusammenhang irgendwie „logisch“ ist. Es gibt immerhin einen gewissen Konsens, der sich aus der Beobachtung der umgebenden Welt ergeben hat. Dazu gehören so einfache Regeln, wie das „dictum de omni“, das „tertium non datur“, usw., aber selten mehr. Dennoch steht damit ein Grundgerüst, an dem sich jeder Kalkül messen lassen muß, der für sich in Anspruch nimmt, das inhaltliche Schließen adäquat abzubilden.

Kalküle sollen dabei helfen, das inhaltliche Denken in die richtigen Bahnen zu lenken, es auch maschinell nachprüfbar zu machen, und damit in Bereiche zu führen, die der menschlichen „Handtaschenlogik“ nur in den seltensten Fällen zugänglich sind.

Eine Vermeidung der Antinomien ist also völlig unzureichend. Eine Lösung ist nötig. Mindestens zwei Möglichkeiten bieten sich an:

1. Es gelingt der Nachweis, daß bei den vermeintlichen Antinomien Regeln des Kalküls falsch angewendet werden, oder die Antinomien auf widersprüchlichen Voraussetzungen beruhen. Dann gibt es keine Antinomien in dem Kalkül.
2. Es gibt die Antinomien in dem betrachteten Kalkül tatsächlich. Dann muß herausgefunden werden, welche Grundformeln oder Grundregeln des Kalküls für die Antinomien verantwortlich sind. Diese müssen dann modifiziert werden. Ist diese Modifikation mit dem inhaltlichen Schließen konform, so besteht ein Unterschied zwischen der bisherigen Formalisierung des inhaltlichen Schließens (dem ursprünglichen Kalkül) und dem inhaltlichen Schließen selbst. Der ursprüngliche Kalkül erfüllt also nicht die an ihn gestellten Bedingungen und muss verbessert werden.

Die in diesem Fall vorzunehmenden Modifikationen sind nicht mit der symptomatischen Behandlung der Antinomien vergleichbar. Dort werden Grundformeln und Grundregeln des Kalküls verändert, weggelassen oder hinzugefügt, von denen nicht nachgewiesen wird, daß sie die Verursacher der Antinomien sind. Diese Veränderungen sind dann oft viel zu weitgehend und verbieten auch durchaus vernünftige Schlüsse, Folgerungen, Konsequenzen, usw.

Widerspricht allerdings selbst die notwendige Modifikation des Kalküls dem inhaltlichen Schließen, so ist die Konsequenz ein totaler Skeptizismus, denn dann ist das inhaltliche Schließen offenbar widersprüchlich und das ganze bisherige wissenschaftliche Denken in einer schweren Krise.

Auf jeden Fall sollte eine Behandlung der Antinomienfrage aber möglichst formal sein, damit Unklarheiten in der Formulierung der Probleme so weit wie möglich ausgeschlossen werden können. Dazu sollen hier Aussagen- und Prädikatenlogik benutzt werden, als die Kalküle, in denen die meisten Antinomien formuliert wurden, und zwar mit den folgenden Formelzeichen:

T	wahr	
F	falsch	
$\neg$	Negation	(Negation auch durch Überstreichen)
$\vee$	oder	
$\wedge$	und	(kann ggf. auch weggelassen werden)
$\Rightarrow$	Implikation	
$\Leftrightarrow$	Äquivalenz	
$\bigwedge$	Allquantor	
$\bigvee$	Existenzquantor	
$,$	Aufzählung	(ähnlich wie $\wedge$)

Die Grundformeln und Schlußregeln der Aussagen- und Prädikatenlogik dürften bekannt sein, daher sollen nur einige wenige wichtige Punkte erwähnt werden:

1.1 Die Mengenlehre

Die Mengenlehre mit ihrer Elementbeziehung $\in$ wird als nicht zur Logik gehörig betrachtet, sondern ist eine auf die Logik aufgesetzte Erweiterung. Wir werden uns aber trotzdem sehr ausführlich mit dem Begriff „Menge“ befassen müssen, denn zahlreiche Antinomien sind von mengentheoretischer Natur.

1.2 Der modus ponens

Der „modus ponens“ sieht formalisiert so aus:

$$A, (A \Rightarrow B) \Rightarrow B$$

Umgangssprachlich heißt das: Ist A wahr und gilt der Ableitungszusammenhang „Wenn A dann B “, so ist auch B wahr, wobei A und B beliebige Aussagen sind. Besonderer Wert ist darauf zu legen, daß gefordert wird, daß A wahr sein muß, damit B gefolgert werden kann. Gelegentlich wird dieser wichtige Aspekt übersehen.

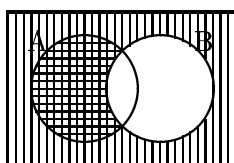
Wie kann man nun „ A muß wahr sein“ verstehen? Es gibt drei Sorten von wohlgeformten Ausdrücken:

1. Die Tautologien, also die immer wahren Ausdrücke. Tautologien kann man ganz bedenkenlos im modus ponens verwenden, sie sind ja immer wahr. Nur solche tautologischen Ausdrücke werden **Sätze** genannt.

2. Die Ausdrücke, die nur unter ganz bestimmten Bedingungen wahr sind. Will man solche Ausdrücke im modus ponens verwenden, so sind diese bestimmten Bedingungen als gegeben anzunehmen, und das muß man sich merken. Es muß vor einer Anwendung des modus ponens also überprüft werden, ob diese Bedingungen überhaupt als „wahr“ angenommen werden können, ob also nicht andere Voraussetzungen dagegen sprechen.
3. Die Negate der Tautologien, die Antilogien. Diese wohlgeformten Ausdrücke sind immer falsch, es gibt keine Bedingungen unter denen diese Ausdrücke jemals wahr sein könnten, weil sie **in sich** widersprüchlich sind. In einem modus ponens sind diese Ausdrücke nicht verwendbar, da die Bedingungen für den modus ponens prinzipiell unerfüllbar wären.

Konsequenz: Bevor man den modus ponens anwendet, muß man sich zunächst darüber im klaren sein, ob dessen Voraussetzungen auch wirklich erfüllt sind. Aus Antilogien, also aus Falschem, kann man bekanntlich alles ableiten, aber das bringt nichts, da man die Ergebnisse nicht als wahr erweisen kann.

Interessanterweise ist die Forderung, daß A wahr sein muß, damit man mit $A \Rightarrow B$ auf die Wahrheit von B schließen kann, eigentlich zu stark. Im Venn-Diagramm wird das deutlich:



Ein schraffierter Bereich markiert die Nicht-Existenz dieses Bereiches. Die horizontale Schraffur gibt den Zusammenhang $A \Rightarrow B$ an, die vertikale Schraffur das Ergebnis des modus ponens, nämlich B . Deutlich ist zu sehen, daß nur noch das Feld $\neg A \wedge \neg B$ schraffiert werden müßte, was der Forderung $\neg A \Rightarrow B$ entspricht, die schwächer ist als A . Es reichen also $A \Rightarrow B$ und $\neg A \Rightarrow B$, um B zu folgern, was einer gewissen Unabhängigkeit des B von A entspricht.

Für den Fall des modus ponens mit zwei Prämissen A und B aus denen die Konklusion C folgt, ergibt sich im Rückschluss auf minimale verborgene Prämissen: $(A, B \Rightarrow C), X \Rightarrow C$ für X : $X = (\neg A \Rightarrow C, \neg B \Rightarrow C)$, statt $X = (A, B)$, wie es der modus ponens verlangt.

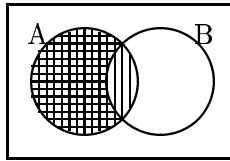
1.3 Der modus tollens

Der „modus tollens“ sieht formalisiert so aus:

$$(A \Rightarrow B), \neg B \Rightarrow \neg A$$

Umgangssprachlich heißt das: Gilt der Zusammenhang „Wenn A dann B “ und ist B nicht wahr, so ist auch A nicht wahr. Wichtig ist hierbei, was geschieht, wenn mehrere ungesicherte Prämissen in den Schluß eingehen. Gilt also z.B.: $A_1, A_2 \Rightarrow B$ und $\neg B$, so folgt daraus: $\neg(A_1 \wedge A_2)$, also: $\neg A_1 \vee \neg A_2$. Der Verursacher des Widerspruches ist demnach nicht eindeutig bestimmbar.

Auch für den modus tollens gilt, daß seine Prämissen eigentlich viel zu stark sind:



Die horizontale Schraffur gibt wieder die Prämisse $A \Rightarrow B$ an, die vertikale Schraffur die Konklusion $\neg A$. Die zusätzliche Bedingung $A \Rightarrow \neg B$ reicht bereits völlig aus, um der Konklusion zu genügen, $\neg B$ ist wieder zu stark. Weiß man statt $A \Rightarrow \neg B$ mehr, nämlich $\neg B$, umso besser.

1.4 Ableiten und Beweisen

Eine Ableitung in einem Kalkül ist eine Folge von Regelnweisungen (Transformationen), die nacheinander angewendet einen Ausdruck des Kalküls von einem Anfangszustand in einen Endzustand überführen. Insofern kann man demonstrieren, daß ein Ableitungszusammenhang besteht, indem die Voraussetzungen jeder Regelnanwendung, sowie die korrekte Anwendung jeder Regelnweisung nachgeprüft und nachgewiesen wird. Allerdings wird die Gültigkeit (unter welchen Bedingungen ist der Ausdruck wahr?) des Anfangsausdruckes nicht überprüft.

Im Gegensatz dazu wollen wir hier eine Ableitung, bei der die Gültigkeit des Anfangsausdruckes nachgewiesen wird einen Beweis nennen. Eine Ableitung ist also etwas Schwächeres als ein Beweis.

1.5 Das Definieren

Entgegen der weit verbreiteten Ansicht, ist auch das Definieren nicht so frei, wie es oft behauptet wird, zumindest wenn man in dem Kalkül auch etwas erreichen will. Das Definieren darf nämlich nicht soweit gehen, daß eine Hinzunahme der Definition zu den logischen Axiomen des (jeweils verwendeten) Kalküls zu einem Widerspruch führt. Dann ist z.B. in der Aussagenlogik (im Rahmen der syntaktischen Möglichkeiten) alles ableitbar, aber die Ergebnisse der Ableitungen sind nicht beweisbar, da die Voraussetzungen für die Beweisbarkeit nicht gegeben sind, d.h. es ist nicht möglich, das Ergebnis der Ableitung von der Ableitung selbst (mit dem modus ponens) zu trennen. Jede Definition muß dahingehend untersucht werden, sonst kann es böse Überraschungen geben. Ausdrücke, von denen man geglaubt hat, sie wären wichtig, können sich beweistechnisch in Nichts auflösen, weil sie dem verwendeten Axiomensystem widersprechen.

2 Der Lügner

Basis dieser Betrachtung des Antinomien-Problems ist der **Lügner** des Epimenides (ca. 600 v.C.), der sich nach [1] so formulieren läßt:

„Was ich jetzt sage, ist falsch“

Nimmt man an, dieser Ausdruck sei wahr, so ergibt sich, daß er falsch ist. Nimmt man andererseits an, er sei falsch, so sagt er aus, daß er wahr ist.

Eine Antinomie? Das Ende der klassischen Logik?

Um das herauszufinden formalisieren wir die umgangssprachliche Formulierung. Sei L das, was der Lügner sagt. Dann gilt:

$L \Leftrightarrow$ „Was ich jetzt sage, ist falsch“

„Was ich jetzt sage, ...“ ist aber gerade L selbst.

Also:

$L \Leftrightarrow (L \text{ ist falsch})$

„ L ist falsch“ kann aussagenlogisch durch „ $L \Leftrightarrow F$ “ ausgedrückt werden, so daß der Lügner die folgende Form annimmt:

$L \Leftrightarrow (L \Leftrightarrow F)$

Die dazugehörige Argumentation zeigt nun folgendes:

$(L \Leftrightarrow (L \Leftrightarrow F), L \Leftrightarrow T) \Rightarrow L \Leftrightarrow F$
 $(L \Leftrightarrow (L \Leftrightarrow F), L \Leftrightarrow F) \Rightarrow L \Leftrightarrow T$

Offenbar wird dabei völlig korrekt geschlossen, nämlich:

$$\frac{\begin{array}{cc} (Lügner) & (Annahme) \\ L \Leftrightarrow (L \Leftrightarrow F) & L \Leftrightarrow T \end{array}}{\frac{T \Leftrightarrow (L \Leftrightarrow F)}{L \Leftrightarrow F}}$$

$$\frac{\begin{array}{cc} (Lügner) & (Annahme) \\ L \Leftrightarrow (L \Leftrightarrow F) & L \Leftrightarrow F \end{array}}{(L \Leftrightarrow F) \Leftrightarrow T}$$

Die übliche Antinomien-Argumentation läßt sich nachvollziehen, es ist also sehr wahrscheinlich (aber das ist hier vorausgesetzt), daß der Lügner richtig formalisiert wurde. Insbesondere wird die Selbstbezüglichkeit der Aussage recht deutlich.

Es wurde dabei nach dem Schlußschema $A, B \Rightarrow C$ geschlossen, also mit zwei Prämissen. Will man nun die Konklusion $C = (L \Leftrightarrow F)$ abtrennen, so müßte man dem modus ponens folgend nachweisen, daß $A = (L \Leftrightarrow (L \Leftrightarrow F))$ und $B = (L \Leftrightarrow T)$ wahr sind, und entsprechend für $C = (L \Leftrightarrow T)$ und $B = (L \Leftrightarrow F)$. Nur dann hätte man aus der „Wahrheit“ $L \Leftrightarrow F$ die „Wahrheit“ $L \Leftrightarrow T$ gewonnen, bzw. umgekehrt aus $L \Leftrightarrow T \quad L \Leftrightarrow F$, was logisch betrachtet recht peinlich wäre. In den Schluß geht aber auch die Definition des Lügners $A = (L \Leftrightarrow (L \Leftrightarrow F))$ mit ein. Das mag unnötig erscheinen, schließlich ist es eine Definition. Eine Definition ist aber nicht deshalb schon in Ordnung, weil es möglich ist, sie aufzuschreiben. Man kann schließlich den größten Unsinn ohne Probleme aufschreiben. Diese Definition geht an zentraler Stelle in den logischen Schluß mit ein und muß daher selbstverständlich auch Objekt logischer Untersuchung sein.

Wie also sieht die Belegung dieser Definition mit den Wahrheitswerten aus? Von der Beantwortung dieser Frage hängt ab, ob der sich ergebende Widerspruch ein relevantes Ergebnis ist:

$$\frac{\text{Setze } L \Leftrightarrow T :}{\frac{T \Leftrightarrow (T \Leftrightarrow F)}{F}}$$

$$\frac{\text{Setze } L \Leftrightarrow F :}{\frac{F \Leftrightarrow (F \Leftrightarrow F)}{F}}$$

Die Voraussetzungen $L \Leftrightarrow (L \Leftrightarrow F)$ (A) und $L \Leftrightarrow T$ bzw. $L \Leftrightarrow F$ (B) widersprechen sich also jeweils, können nicht zusammen wahr sein, die Definition des Lügners A erweist sich sogar als eine Antilogie, das Gegenteil einer Tautologie, also ein immer falscher Ausdruck. Mit dem modus ponens ist daher die Konklusion $L \Leftrightarrow F$ bzw. $L \Leftrightarrow T$ (C) nicht abzutrennen, denn A kann nicht als wahr erwiesen werden.

Das Ergebnis des modus ponens gilt aber, wie bereits oben bemerkt, auch schon unter schwächeren Bedingungen. Es reicht nachzuweisen, daß gilt: $\neg A \Rightarrow C$ und $\neg B \Rightarrow C$. A ist eine Antilogie, $\neg A$ daher eine Tautologie, daher kann aus $\neg A$ nicht $L \Leftrightarrow F$ oder $L \Leftrightarrow T$ folgen. D.h. auch die minimale Schlußbedingung ist nicht erfüllbar. C kann nicht abgetrennt werden.

Konsequenz: Erweist sich eine Definition (ein Urteil) als in sich widersprüchlich, so kann eine Konklusion, die auf dieser Voraussetzung beruht weder nach den modus ponens, noch nach der minimalen Abtrennungsbedingung abgetrennt werden. Und bei den Antinomien ist die Konklusion in keinem Fall tautologisch, wäre dem so, so wäre der Sinn der Antinomien verfehlt.

Anwendbar ist in diesem Fall der modus tollens, der das folgende Ergebnis liefert:

Voraussetzung: $L \Leftrightarrow (L \Leftrightarrow F) \Rightarrow L \Leftrightarrow \neg L$ da natürlich $(L \Leftrightarrow F) \Leftrightarrow \neg L$ gilt

Es gilt aber: $\neg(L \Leftrightarrow \neg L)$ $L \Leftrightarrow \neg L$ ist immer falsch

Daraus folgt: $\neg(L \Leftrightarrow (L \Leftrightarrow F))$

D.h., wie oben ergibt sich, daß der Lügner immer falsch ist. Eine dritte Möglichkeit, dieses zu zeigen ist, aus Tautologien das Negat des Lügners herzuleiten:

$$\frac{\frac{\frac{(Tautologie)}{L \Leftrightarrow (L \Leftrightarrow T)}}{\neg L \Leftrightarrow \neg(L \Leftrightarrow T)}}{\neg L \Leftrightarrow (L \Leftrightarrow F)} \quad \frac{(Tautologie)}{\neg(L \Leftrightarrow \neg L)}$$

$$\frac{\neg(L \Leftrightarrow (L \Leftrightarrow F))}{\neg(L \Leftrightarrow (L \Leftrightarrow F))}$$

Der Lügner ist keine Antinomie, denn es wird nicht aus widerspruchsfreien Prämissen geschlossen. Ein Ausdruck wie „Was ich jetzt sage, ist falsch“ mag noch so einfach zu formulieren sein, logisch betrachtet ist er völliger Unsinn und hat in logischen Schlüssen nichts anderes als Unsinn zur Folge. Die Behauptung, der Lügner sei eine Antinomie, beruht auf einer ungenügenden Reflektion über die Folgen von Widersprüchen, sowie auf einer naiven Einstellung zur Widerspruchsfreiheit von Definitionen.

3 Einige neuzeitliche Antinomien

Die nun folgenden Antinomien sind im Zuge der Einführung der Mengenlehre entstanden. Cantor entdeckte erste Antinomien 1895 und 1899, nahm sie jedoch nicht ernst. Erst Russell trat mit einer mengentheoretischen Antinomie 1902 an die Öffentlichkeit (Brief an Frege.). In der Folgezeit entstanden sehr schnell eine ganze Reihe weiterer Antinomien. Erwähnt seien hier noch ohne Anspruch auf Vollständigkeit die Namen: König, Berry, Finsler, Grelling, Nelson, Bartlett, Richard, Curry, Burali-Forti, Quine, Montague und Yuting. Eine umfangreiche Zusammenstellung findet sich in [3].

3.1 Russell

Das erste Untersuchungsobjekt ist die Antinomie von **Russell** (1902). Die Darstellung stammt aus [1]:

Russell bildet die Menge aller der Dinge, die sich nicht selbst als Element enthalten. Ein „ x “ gehört zu dieser Menge – sie heiße „ r “ – genau dann, wenn dieses „ x “ sich nicht selbst als Element enthält. Formalisiert heißt das:

$$\bigwedge_x (x \in r \Leftrightarrow \neg(x \in x))$$

Hieraus läßt sich sehr schnell ein Widerspruch herleiten:

- | | | |
|-----|---|--|
| (1) | $\bigwedge_x (x \in r \Leftrightarrow \neg(x \in x)) \Rightarrow (r \in r \Leftrightarrow \neg(r \in r))$ | spezialisiere x nach r |
| (2) | $\bigwedge_x (x \in r \Leftrightarrow \neg(x \in x))$ | Das ist die Definition der Menge,
Russell setzt sie als wahr an |
| (3) | $r \in r \Leftrightarrow \neg(r \in r)$ | modus ponens aus (1) und (2) |

Nimmt man an, „ r “ gehöre zu der Menge der Dinge, die sich nicht selbst als Element enthalten, so ergibt sich, daß „ r “ nicht dazugehört und umgekehrt. Die Frage, ob die Russellmenge sich selbst enthält, führt also auf einen Widerspruch. Aber dieser Schluß ist nur relevant, wenn nachgewiesen werden kann, daß sich die widersprüchliche Konklusion abtrennen läßt, d.h. also, ob Russell die Definition seiner Menge wirklich als uneingeschränkt wahr ansetzen durfte.

Zur Beantwortung dieser Frage dient die nun folgende Herleitung eines prädikatenlogischen Satzes für eine beliebige zweistellige Relation:

- | | | |
|-----|---|---|
| (1) | $\bigwedge_x (A(x, y) \Leftrightarrow \neg A(x, x)) \Rightarrow (A(y, y) \Leftrightarrow \neg A(y, y))$ | spezialisiere x nach y |
| (2) | $\neg(A(y, y) \Leftrightarrow \neg A(y, y))$ | $A(y, y) \Leftrightarrow \neg A(y, y)$ ist ein
Widerspruch |
| (3) | $\neg \bigwedge_x (A(x, y) \Leftrightarrow \neg A(x, x))$ | modus tollens aus (1) und (2) |

Wird die beliebige zweistellige Relation z.B. durch die Elementbeziehung ersetzt (bedeutet also $A(x, y)$ nun $x \in y$), so ergibt sich:

$$\neg \bigwedge_x (x \in r \Leftrightarrow \neg(x \in x))$$

In Worten: Nicht für alle „ x “ gilt: „ x “ ist Element von „ r “ genau dann, wenn „ x “ nicht Element von sich selbst ist.

Russells Definition der Menge „ r “ behauptet aber das Gegenteil, nämlich: Für alle „ x “ gilt: „ x “ ist Element von „ r “ genau dann, wenn „ x “ nicht Element von sich selbst ist.

Das ist das Negat des oben abgeleiteten Satzes. D.h., Russell setzt einfach voraus, daß sich diese Menge „ r “ bilden läßt, was aber schlicht als falsch bezeichnet werden muß. Die Bedingung der Russellmenge, für alle „ x “ zu gelten, ist aus rein logischen Gründen nicht erfüllbar, ihre Definition ist widersprüchlich, und damit muß die Herleitung des Widerspruches durch Spezialisierung und modus ponens zurückgewiesen werden.

Auch die minimale Abtrennungsbedingung ist nicht erfüllbar, da die Definition der Russellmenge eine Antilogie ist, deren Negat somit eine Tautologie, aus der die nicht-tautologische Konklusion nicht folgen kann.

3.2 Einschub: Der Begriff „Menge“

Diese erste mengentheoretische Antinomie soll Anlaß sein, sich einmal genauer mit dem Begriff „Menge“ zu beschäftigen.

Das Definieren und Bilden von Mengen ist offensichtlich nicht ganz so einfach, wie man vielleicht glaubt. Die Logik darf man dabei nicht vergessen. Anscheinend kann man nicht zu jeder beliebigen Eigenschaft ein Mengen-Individuum bilden, das genau die unter die betreffende definierte Eigenschaft fallenden Individuen als Elemente enthält.

Die Frage ist also: Wie kommt es, oder wie ist es möglich, daß es Dinge gibt, die nicht zu einer Menge zusammengefaßt werden können?

Paul Finsler schreibt dazu in [2], im Aufsatz: „Gibt es Widersprüche in der Mathematik?“, Seite 151 – 153:

...

Wir müssen uns fragen: Was ist eine Menge? Nach Cantor ist es „die Zusammenfassung wohlbestimmter Objekte zu einem Ganzen“. Also die Zusammenfassung selbst. Und es scheint so, als ob man irgendwelche Dinge stets zu einem Ganzen zusammenfassen könnte. Was sollte einen auch daran hindern? Und doch muß man hier vorsichtig sein. Ja, wenn es sich um konkrete, etwa materielle Dinge handelt, oder überhaupt um solche, die an sich mit der Bildung von Mengen gar nichts zu tun haben, dann liegt in der Tat keinerlei Grund vor, warum man sie nicht sollte zusammenfassen können.

Anders ist es aber, wenn die Dinge oder ihre Zugehörigkeit zur Menge in Abhängigkeit treten kann von der zu bildenden Menge selbst. Dann kann diese Abhängigkeit von solcher Art sein, daß sie nicht erfüllt werden kann.

Genauer ist zu sagen: Wenn man auch Mengen von Mengen bildet, oder allgemein Dinge als Elemente zuläßt, die von Mengen abhängig sind, dann ist die Definition Cantors eine *Zirkeldefinition*, und eine solche braucht nicht stets erfüllbar sein.

Dieses letztere möchte ich, da es wichtig ist, an einem einfachen Beispiel erläutern.

Man kann eine Zahl x durch einen algebraischen Ausdruck definieren, etwa durch den Ausdruck $a - b$. Es soll also

$$x = a - b$$

sein. Wenn nun a und b feste Zahlen sind, die nicht von x abhängen, so liegt kein Zirkel vor, die Definition ist stets erfüllbar, es gibt stets eine Zahl x , die ihr genügt.

Nun kann aber der algebraische Ausdruck, durch den x definiert werden soll, von x selbst abhängen. Das ist dann eine Zirkeldefinition und eine solche kann unter Umständen ebenfalls eindeutig erfüllbar sein. Wenn etwa

$$x = a - x$$

sein soll, so genügt dieser Forderung die Zahl $x = \frac{a}{2}$ und nur diese.

Eine Zirkeldefinition braucht aber nicht erfüllbar sein; bei

$$x = a + x$$

z.B. gibt es, wenn a von Null verschieden ist, keine Zahl, die ihr genügt.

Man kann die Zirkel nicht einfach verbieten, solche Gleichungen kommen selbst in der angewandten Mathematik oft genug vor, sondern man muß mit ihnen zu rechnen wissen. Genauso ist es bei Mengen: Cantors allgemeine Definition ist eine Zirkeldefinition, da darf es einen nicht wundern, daß es Fälle gibt, in denen sie versagt.

...

Es ist noch ein Einwand zu besprechen. Man sagt etwa:

Wenn man sich alle wirklich existierenden Mengen, die sich nicht selbst enthalten, gegeben denkt, so kann man sie sich doch *vorstellen*, daß nun diese alle wieder zu einem Ganzen, zu einer Menge zusammengefaßt werden.

Darauf ist zu erwidern: Das ist nicht wahr! Das kann man sich nicht vorstellen. Man stellt sich irgend etwas anderes dabei vor, aber nicht das worauf es ankommt. Denn, wenn man sich dies vorstellen könnte, so müßte man doch auch sagen können, ob man diese Zusammenfassung selbst mit zusammenfaßt oder nicht, ob man sich diese Menge als eine solche vorstellt, die sich enthält, oder als eine solche, die sich nicht enthält; und das kann man nicht sagen. Nein, etwas logisch Widerspruchsvolles kann man sich überhaupt niemals vorstellen, diese Menge ebensowenig wie etwa die größte natürliche Zahl, die es eben auch nicht gibt.

...

Soweit Paul Finsler. Durch Georg Cantor wurde der Begriff der „Menge“ in die Mathematik eingeführt. Er definiert 1895 in [8] den Begriff „Menge“ etwas anders, nämlich so:

...

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

...

Das ist ein kleiner, zumindest verbaler Unterschied zu Finslers Version der Cantorsche Definition. Cantor selbst hat die Definition der Menge im Laufe der Zeit immer genauer gefaßt. In den frühen Arbeiten (1872 – 1880) noch sehr undeutlich definiert, schreibt Cantor in [4] aus dem Jahre 1882 über Mengen:

...

Eine Mannigfaltigkeit (ein Inbegriff, eine Menge) von Elementen, die irgendwelcher Begriffssphäre angehören, nenne ich *wohldefiniert*, wenn auf Grund ihrer Definition und infolge des logischen Prinzips vom ausgeschlossenen Dritten es als *intern bestimmt* angesehen werden muß, *sowohl* ob irgendein derselben Begriffssphäre angehöriges Objekt zu der gedachten Mannigfaltigkeit als Element gehört oder nicht, *wie auch*, ob zwei zur Menge gehörige Objekte, trotz formaler Unterschiede in der Art des Gegebenseins einander gleich sind oder nicht.

...

In [5] aus dem Jahre 1883 schreibt Cantor dann in den Anmerkungen:

...

1) (S. 545.) *Mannigfaltigkeitslehre*. Mit diesem Worte bezeichne ich einen sehr viel umfassenderen Lehrbegriff, den ich bisher nur in der speziellen Gestaltung einer arithmetischen oder geometrischen Mengenlehre auszubilden versucht habe. Unter einer „Mannigfaltigkeit“ oder „Menge“ verstehe ich nämlich allgemein jedes Viele, welches sich als Eines denken läßt, d.h. jeden Inbegriff bestimmter Elemente, welcher durch ein Gesetz zu einem Ganzen verbunden werden kann, ...

In [6] aus dem Jahre 1887 schreibt er:

...

Unter *Mächtigkeit* oder *Kardinalzahl* einer Menge M (die aus wohlunterschiedenen, begrifflich getrennten Elementen $m, m', \dots$ besteht und insofern bestimmt und abgegrenzt ist) verstehe ich ...

Es folgt die bereits oben zitierte Definition aus dem Jahre 1895 und dann schreibt Cantor (siehe [9], Seite 443) in einem Brief an Dedekind vom 28. Juli 1899:

...

Gehen wir von dem Begriff einer bestimmten Vielheit (eines Systems, eines Inbegriffes) von Dingen aus, so hat sich mir die Notwendigkeit herausgestellt, zweierlei Vielheiten (ich meine *bestimmte* Vielheiten) zu unterscheiden.

Eine Vielheit kann nämlich so beschaffen sein, daß die Annahme eines „Zusammenseins“ *aller* ihrer Elemente auf einen Widerspruch führt, so daß es unmöglich ist, die Vielheit

als eine Einheit, als „ein fertiges Ding“ aufzufassen. Solche Vielheiten nenne ich *absolut unendliche* oder *inkonsistente Vielheiten*.

...

Wenn hingegen die Gesamtheit der Elemente einer Vielheit ohne Widerspruch als „zusammenseiend“ gedacht werden kann, so daß ihr Zusammengefaßtwerden zu „*einem* Ding“ möglich ist, nenne ich sie eine *konsistente Vielheit* oder eine „Menge“.

...

Das ist nun fast genau Finslers Version der Cantorschen Mengendefinition. Daher kann man die Frage nach den nicht-wohlbestimmten Objekten eindeutig klären. Die Frage, ob solche Objekte, wenn man überhaupt von Objekten sprechen kann, zu Mengen zusammengefaßt werden können, muß, wenn man Finslers Version oder der späten Cantorschen Mengendefinition folgt, mit einem eindeutigen „Nein“ beantwortet werden.

Und in den frühen Mengendefinitionen ist zumindest keine Rede davon, daß es für eine beliebige Eigenschaft immer die Menge der Dinge geben **muß**, auf die diese Eigenschaft zutrifft. Immerhin soll eine Menge „Eines“, also auch widerspruchsfrei sein. Das muß man ersteinmal für die zu definierende Menge sicherstellen, selbstverständlich ist das nicht. Daß man Beliebiges immer zu einer Menge zusammenfassen kann, ist eine Überinterpretation, auch schon der ersten „naiven“ Definitionen des Begriffes „Menge“.

Die Russellsche Antinomie beruht auf eben dieser Vorstellung, daß es zu jeder Zusammenfassung von „Dingen“ ein Mengenindividuum gibt, das genau die unter diese Eigenschaft fallenden Individuen als Elemente enthält.

Ein bißchen Logik, und man sieht, daß man von falschen Voraussetzungen ausgegangen ist. Und damit ist das Problem dann auch erledigt. Änderungen an der Logik sind jedenfalls nicht notwendig und anscheinend auch nicht am naiven Mengenbegriff Cantors. Cantor kann man höchstens vorwerfen, daß die Definition des Begriffes „Menge“ zunächst etwas unscharf war, was aber beim Aufbau eines neuen Zweiges der Mathematik nur zu verständlich ist.

Die Antinomien von Russell und anderen waren aber der Auslöser für die Axiomatisierung der Mengenlehre. Cantors Mengenbegriff wurde als unhaltbar verworfen, und mit Hilfe der axiomatischen Methode durch ein System ersetzt, das die wesentlichen Ergebnisse der Mengenlehre herzuleiten gestattete, die bekannten Antinomien jedoch vermied. Mit der Axiomatisierung der Mengenlehre begann Ernst Zermelo, dessen Arbeit von Adolf Fraenkel fortgesetzt wurde. Auch A. Schoenflies darf in diesem Zusammenhang nicht unerwähnt bleiben. Kennzeichnend für die Axiomatiken vom Zermelo-Fraenkel Typ ist das Aussonderungs(Komprehensions)-Axiom, welches die Bildung von Mengen gestattet und Cantors Mengendefinition praktisch ersetzt. Es besitzt jedoch einige Einschränkungen, die die Antinomien wirkungsvoll verhindern. John v. Neumann verbesserte und vereinfachte das Zermelo-Fraenkelsche Axiomensystem, Paul Bernays und Kurt Gödel fügten den Klassenbegriff der Mengenlehre hinzu und bildeten die Bernays-Gödel Axiomatik der Mengenlehre.

So gut die Axiomatisierung der Mengenlehre getan hat, so wenig gegeben ist aber deren Motivation, die Unhaltbarkeit des Cantorschen Mengenbegriffes. Das eigentliche Problem hat überhaupt nichts mit der Mengenlehre oder dem Mengenbegriff zu tun, sondern ist eine Frage der rein logischen Behandlung widersprüchlicher Definitionen. Es scheint also sinnvoll, nicht von einem „naiven Mengenbegriff“ zu sprechen, sondern wohl eher von den „naiven Anwendern der Mengenlehre“.

Daß es schwierig sein kann Widersprüche aufzufinden, bestreitet niemand. So sieht man einer mit üblichen Zeichen aufgeschriebenen Version der Definition der Russellmenge $r = \{x | x \notin x\}$ ihre Widersprüchlichkeit zunächst nicht an. Was soll auch gefährlich sein an der Eigenschaft $x \notin x$? Schreibt man die Definition aber ausführlich hin, beachtet also den enthaltenen Allquantor, sowie die Tatsache, daß die Mengendefinition ebenfalls eine Elementbeziehung impliziert, so ergibt sich, daß die Russellmenge nicht existiert, da sie so zumindest widersprüchlich definiert ist.

Eine hervorragende Beschreibung der zugrundeliegenden Sachverhalte findet sich in [2] in den Aufsätzen „Über die Grundlegung der Mengenlehre“ Teil 1 und 2.

Für das Verständnis der noch folgenden Abhandlungen, insbesondere zu den Diagonalverfahren, sind Finslers obige Ausführungen über Zirkeldefinitionen sehr wichtig.

3.3 Grelling

Hinter der Antinomie von **Grelling** (1908) steckt nach [1] das gleiche Schema wie hinter Russells Konstruktion. Nur sieht sie mehr „semantisch“ aus, weil zu ihrer Formulierung die „Bedeutungsrelation“ benutzt wird:

Bei den Adjektiven einer Sprache kann man fragen, ob die Bedeutung eines bestimmten Adjektives auf dieses Adjektiv selbst zutrifft. So ist z.B. das Adjektiv „dreisilbig“ selbst dreisilbig, „einsilbig“ ist dagegen nicht einsilbig. Ein Adjektiv auf das seine eigene Bedeutung zutrifft, heißt autolog, sonst heißt es heterolog. Führt man „ $B(x, y)$ “ als Abkürzung ein für: ‚Die Bedeutung des Adjektives „ y “ trifft auf das Adjektiv „ x “ zu‘ so kann man „heterolog“ (h) definieren durch:

$$\bigwedge_x (B(x, h) \Leftrightarrow \neg B(x, x))$$

In Worten: Für alle Adjektive „ x “ gilt: Die Bedeutung des Adjektives „heterolog“ trifft auf das Adjektiv „ x “ genau dann zu, wenn die Bedeutung des Adjektives „ x “ nicht auf das Adjektiv „ x “ zutrifft.

Daraus folgt analog zu Russells Argumentation der Widerspruch durch Spezialisierung nach „ h “:

$$B(h, h) \Leftrightarrow \neg B(h, h)$$

Die Frage also, ob „heterolog“ heterolog ist, führt auf einen Widerspruch. Ebenfalls analog zur obigen Argumentation läßt sich aber wieder mit Hilfe des modus tollens das Negat der Definition von „heterolog“ herleiten. D.h., die obige Definition gilt aus rein logischen Gründen nicht, wie sie selbst behauptet, für alle Wörter, nämlich nicht für „heterolog“ selbst. Die angebliche Antinomie beruht wiederum auf einer widersprüchlichen Prämisse. Die Eigenschaft „heterolog“ ist einfach falsch (widersprüchlich) definiert.

Und die Eigenschaft „autolog“ (a) ist zwar logisch einwandfrei definiert ($\bigwedge_x (B(x, a) \Leftrightarrow B(x, x))$), sagt aber nichts über das Wort „autolog“ selbst aus, denn dafür behauptet es die Trivialaussage: $B(a, a) \Leftrightarrow B(a, a)$ (und insofern stellt sich die Frage, ob es sich überhaupt um eine brauchbare Definition handelt).

3.4 Richard

Erwähnenswert ist noch die Antinomie von **Richard** (1905):

Wolfgang Stegmüller schreibt dazu in [11], Seite 3 – 4:

...

Für die Bildung der Antinomie von Richard betrachten wir die Ausdrücke der deutschen Sprache, welche Definitionen von Eigenschaften natürlicher Zahlen darstellen (wir wollen im folgenden statt „natürliche Zahl“ einfach „Zahl“ sagen). Da die Anzahl der Ausdrücke, welche wir in einer Sprache bilden können, abzählbar ist, muß insbesondere die Klasse jener Definitionsausdrücke abzählbar sein. Wir können diese Ausdrücke somit numerieren und als eine unendliche Folge anschreiben:

$$(a) \quad A_1, A_2, A_3, \dots$$

Die Anordnung kann ganz willkürlich vorgenommen werden. Man kann z.B. bestimmen, daß ein A_i einem A_j dann voranzugehen habe, wenn A_i weniger Buchstaben enthält als A_j , oder, falls beide dieselbe Anzahl von Buchstaben besitzen, dann, wenn der erste unter den vom Beginn des Ausdruckes an gezählten Buchstaben von A_i , der dem entsprechenden Buchstaben in A_j verschieden ist, im Alphabet an früherer Stelle steht als der entsprechende Buchstabe in A_j (lexikographische Anordnung). Da es sich bei all diesen Prädikaten A_i um Zahlprädikate handelt, muß, wenn irgendein A_i herausgegriffen wird, für jede beliebige Zahl entweder gelten, daß diese Zahl die durch jenes A_i bezeichnete Eigenschaft entweder besitzt oder daß sie diese Eigenschaft nicht besitzt. Da die A_i durch ihre Indizes numeriert werden, kann man dieses auch so ausdrücken: Für zwei beliebig herausgegriffene Zahlen n und k muß entweder der Fall eintreten, daß n die durch A_k bezeichnete Eigenschaft besitzt oder daß n die durch A_k bezeichnete Eigenschaft nicht besitzt. Ist der erste Fall gegeben, so schreiben wir abkürzend „ $A_k(n)$ “, während wir für den zweiten Fall die Abkürzung „ $\sim A_k(n)$ “ benützen. Wir betrachten nun die Eigenschaft, welche mittels der Formel „ $\sim A_n(n)$ “ (1) ausgedrückt wird. Dies ist offenbar eine in der deutschen Sprache definierte Eigenschaft; denn diese Formel besagt ja: „ n hat nicht die Eigenschaft, welche durch A_n bezeichnet wird“ (2), und da laut Voraussetzung A_n ein Ausdruck der deutschen Sprache ist, so gilt dies auch von Satz (2) für den die Formel (1) nur eine Abkürzung darstellt. Die durch (1) bzw. (2) definierte Eigenschaft muß somit, da die Folge (a) alle deutschen Ausdrücke enthält, welche Zahleigenschaften definieren, mit einem dieser A_i zusammenfallen, d.h. es muß eine Zahl r geben, so daß für jede beliebige Zahl n die beiden Bedingungen $A_r(n)$ und $\sim A_n(n)$ zusammenfallen. Was für beliebiges n gilt, muß insbesondere für die spezielle Zahl r gelten. Es müßte also $A_r(r)$ dasselbe sein, wie $\sim A_r(r)$. Dies ist offenbar ein Widerspruch, da die zweite Formel gerade die Negation der ersten darstellt.

...

D.h. also: Für alle n gilt: n besitzt die durch „ A_r “ bezeichnete Eigenschaft fällt zusammen mit n besitzt **nicht** die durch „ A_n “ bezeichnete Eigenschaft. In Formeln:

- (1) $\bigwedge_n (A_r(n) \Leftrightarrow \sim A_n(n)) \Rightarrow A_r(r) \Leftrightarrow \sim A_r(r)$ Richards Argumentation
- (2) $\bigwedge_n (A_r(n) \Leftrightarrow \sim A_n(n))$ Richards Definition, die als wahr angenommen wird
- (3) $A_r(r) \Leftrightarrow \sim A_r(r)$ modus ponens aus (1) und (2)

Das ist natürlich ein Widerspruch.

Ein Beispiel für diese Antinomie ist der Ausdruck: „Die kleinste mit weniger als 100 Buchstaben nicht angebbare Zahl“. Die Zahl der Buchstaben dieser Definition ist kleiner als 100. Man kann die Zahl also doch mit weniger als 100 Buchstaben definieren.

Eine Antinomie?

Zusammengefaßt heißt das obige Beispiel aber nichts anderes als: „Die kleinste mit weniger als 100 Buchstaben nicht angebbare Zahl ist mit weniger als 100 Buchstaben angebbar“.

Das ist ein widersprüchlicher Zahlenausdruck. Die Frage ist: Dürfen „widersprüchlich“ definierte Ausdrücke in der von Richard aufgestellten Liste vorkommen?

Er schreibt selbst: „Für zwei beliebig herausgegriffene Zahlen n und k muß **entweder** der Fall eintreten, daß n die durch A_k bezeichnete Eigenschaft besitzt **oder** daß n die durch A_k bezeichnete Eigenschaft **nicht** besitzt.“ Das gilt nur für widerspruchsfrei definierte Ausdrücke, ein widersprüchlicher Ausdruck ist gerade dadurch charakterisiert, daß ihm irgendeine Eigenschaft sowohl zukommt, als auch nicht zukommt. Das macht gerade seine Widersprüchlichkeit aus. Die Liste darf also nach Richards eigenen Vorgaben nur widerspruchsfrei definierte Zahlausdrücke enthalten. Und nur dann kann er ja auch mit Hilfe des modus ponens oder der minimalen Abtrennungsbedingung etwas daraus folgern.

Die von Richard definierte Zahleigenschaft ist aber widersprüchlich, denn für jedes n ist die Definition $\bigwedge_n (A_r(n) \Leftrightarrow \sim A_n(n))$, analog zur Argumentation bei Russell und Grelling, in mindestens einem Fall aus rein logischen Gründen nicht erfüllbar, d.h. der modus ponens und auch die minimale Abtrennungsbedingung sind nicht anwendbar, und es gibt entgegen Richard's Behauptung **keine** Zahl r , so daß für **jedes** n $A_r(n)$ und $\sim A_n(n)$ zusammenfallen. Das ist auch so in Ordnung, denn widersprüchlich definierte Zahlausdrücke kommen ja in der Liste, nach Richards eigener Vorgabe, nicht vor.

Und würde man dieses erlauben, so wäre der ableitbare Widerspruch $A_r(r) \Leftrightarrow \sim A_r(r)$ nicht verwunderlich, man hat ihn ja geradezu herausgefordert, indem man Ausdrücke in der Liste zugelassen hat, für die sowohl $A_r(r)$ als auch $\sim A_r(r)$ gelten kann.

Richard scheint davon auszugehen, daß jeder Zahlausdruck, nur weil er in der deutschen Sprache ausdrückbar (aufschreibbar) ist, automatisch auch widerspruchsfrei ist. Nur so kann man sein bedenkenloses Vorgehen erklären. Auch hier zeigt sich wieder: Die Antinomie ist keine, denn der Widerspruch, der am Ende herauskommt, ist bereits als Prämisse hineingesteckt worden, offenbar entstanden aus einem unreflektierten Umgang mit der Frage der Widerspruchsfreiheit von Definitionen.

Gelegentlich findet man Ablehnungen der Richardschen Antinomie (z.B. in [13]), die darauf beruhen, daß die betrachteten Ausdrücke nur rein arithmetische Eigenschaften sein sollen. Die Richardsche Definition ist keine rein arithmetische Eigenschaft, sondern mit umgangssprachlichen Wörtern durchsetzt, daher tritt der Widerspruch nicht auf. Das ist keine Lösung, da die Voraussetzungen einseitig so verändert wurden, daß der Widerspruch nicht auftritt (ein typisch symptomatisches Vorgehen), vor allem wird damit nicht das von Stegmüller vorgelegte Problem gelöst.

Wie oben zu sehen war, ist ein solches Vorgehen gar nicht nötig, auch wenn man z.B. die Richardsche Definition rein arithmetisch darstellen könnte, oder aber zuläßt, daß auch nicht rein arithmetische Eigenschaften in der Liste vorkommen können. Auch dann tritt der Widerspruch nicht wirklich auf, weil die Richardsche Eigenschaft nicht wohldefiniert ist, und daher der Widerspruch gar nicht beweisbar ist.

4 Allgemeine Bemerkungen

Nach diesen Beispielen ist es nun an der Zeit, sich einmal allgemeinere Gedanken über das Antinomien-Problem zu machen.

4.1 Die Typentheorie

Eine sehr beliebte Methode, gegen die Antinomien anzugehen, ist die Schichten- oder auch Typentheorie. Dabei geht man davon aus, daß Ausdrücke und Ausdrücke über Ausdrücke verschiedenen Schichten angehören, die nicht zirkulär miteinander vermengt werden dürfen. In der Antinomie von Richard ist die Eigenschaft, daß n nicht die Eigenschaft hat, die durch A_n bezeichnet wird, ein solcher Ausdruck einer höheren Schicht. Auch Russells Menge der Dinge, die sich nicht selbst als Element enthalten, ist eine Mengenbeschreibung von Mengen, gehört also einer höheren Schicht an. In der Tat verhindert diese Methode (strenge Typentrennung) das Auftreten der Antinomien, aber sie löst diese Antinomien nicht, erklärt nicht ihre Ursache. Gut erkennbar ist das an Grellings Antinomie. Bei dieser Antinomie wäre man nach der Typentheorie gezwungen, verschiedene Sprachschichten anzunehmen, die nicht miteinander vermengt werden dürfen. Es ist aber die Frage, worin sich eine Sprache und eine Sprache über diese Sprache unbedingt unterscheiden müssen, und warum diese nicht vermengt werden dürfen. Es ist kein Problem, in einer Sprache über diese Sprache selbst zu sprechen, wenn sie nur genügend mächtig ist, und es ist nicht einzusehen, warum dieses grundsätzlich verboten sein sollte. Zirkel können unerfüllbar sein, sie müssen es aber nicht. Daher geht ein generelles Verbot von schichtenübergreifenden Zirkeln viel zu weit, es wird lediglich an Symptomen kurriert. Außerdem begibt man sich in die Gefahr, aus Versehen wichtige und nützliche Dinge zu verbieten.

4.2 Eine allgemeine Lösung für Antinomien

Nun wollen wir ein allgemeines Schema für die Lösung von Antinomien angeben:

Gehorcht eine Argumentation dem Schema:

Aus A und der Definition von A folgt $\neg A$
 und: Aus $\neg A$ und der Definition von A folgt A

so läßt sich ganz formal ableiten, daß die Definition von A logisch falsch ist, und zwar ohne dabei an irgendeiner Stelle auf die Widerspruchsfreiheit der Logik Bezug nehmen zu müssen. Dabei sei X die Formalisierung der Definition von A . Wie X im Detail aussieht, ist vollkommen unerheblich:

$$\frac{\frac{A, X \Rightarrow \neg A}{X \Rightarrow \neg A} \quad \frac{\neg A, X \Rightarrow A}{X \Rightarrow A}}{X \Rightarrow \neg A \wedge A} \\ \frac{X \Rightarrow \neg A \wedge A}{X \Rightarrow F} \\ X \Leftrightarrow F$$

X , also die Definition von A , ist falsch, darf also in einem modus ponens nicht als Voraussetzung verwendet werden. Außerdem ist auch die minimale Abtrennungsbedingung nicht anwendbar, denn aus dem Negat der falschen Definition kann der nicht-tautologische Widerspruch nicht folgen. Behauptet man trotzdem, daß X gilt, so hat man den Widerspruch, den man am Ende einer antinomischen Argumentation erhält, bereits mit der Definition von A hineingesteckt.

Anders ausgedrückt:

Auf rein logischer Basis können die „Antinomien“ beigelegt werden, indem ihre widersprüchlichen Voraussetzungen aufgedeckt werden. Das läßt sich ganz allgemein durchführen, denn:

Folgt aus den Prämissen $P_1, P_2, P_3, \dots, P_n$ ein Widerspruch, gilt also:

$$P_1, P_2, P_3, \dots, P_n \Rightarrow F$$

so folgt nach dem modus tollens: $\neg(P_1, P_2, P_3, \dots, P_n)$, d.h. mindestens eine der Prämissen widerspricht mindestens einer anderen Prämisse, oder aber logischen Axiomen.

Wie kann es sein, daß diese einfachen Schlüsse nicht beachtet werden? D.h., wie kann es sein, daß erstens der zur Gewinnung eines vernünftigen Ergebnisses für den modus ponens mindestens notwendige Widerspruchsfreiheitsnachweis der Prämissen einfach „vergessen“, oder nicht vollständig durchgeführt wird und daß zweitens der modus tollens als Lösungsmöglichkeit gelegentlich vollständig „vergessen“ wird.

Das Problem liegt offensichtlich in der Art und Weise in der definiert bzw. mit Definitionen gearbeitet wird. Vorstellbar ist, daß die „Freiheit des Definierens“ zu wörtlich genommen wird. Es steht natürlich frei, beliebige Ausdrücke zu definieren, die den syntaktischen Regeln des jeweils verwendeten Kalküls entsprechen, aber nicht alle diese Definitionen sind automatisch auch widerspruchsfrei, ganz im Gegenteil. Das wäre so, als ob man verlangen würde, daß es in der Deutschen Sprache unmöglich sei, etwas Widersprüchliches zu sagen oder aufzuschreiben. Diese Meinung kann niemand ernstlich vertreten.

Das Definieren ist logisch gesehen etwas sehr starkes, nämlich quasi das Hinzufügen weiterer Axiome zu den bisherigen Axiomen des verwendeten Kalküls.

Das ist zu unterscheiden von einer reinen Nominaldefinition, die einer bereits bekannten Sache einen neuen Namen gibt. Bei Definitionen steckt oft mehr dahinter, und das macht dann logische Untersuchungen nötig.

4.3 Antinomien?

Es gibt, so scheint es zumindest, keine Antinomien in der klassischen Logik. Der Glaube, es gäbe Antinomien, liegt offenbar darin begründet, daß man Antilogien als Voraussetzungen verwendet, ohne diese als solche zu erkennen (oder erkennen zu wollen?).

Warum werden einfachste logische Regeln nicht beachtet, besonders wenn sich ein Widerspruch ergibt? Es läßt sich aus einem Ausdruck ein Widerspruch herleiten, aber der Ausdruck wird nicht unter die Lupe genommen, es wird nicht überprüft, ob dieser Ausdruck bereits widersprüchlich ist, es wird nicht nach verborgenen Prämissen gesucht, sondern die Logik wird des Versagens bezichtigt.

Es ist zwar ein gewisses Defizit an logischem Verständnis in Teilen der großen Gemeinde der Mathematiker zu beobachten (Logik wird gelegentlich als trivial und primitiv abgetan!), aber daß dieses Defizit schon mit dem *modus ponens* und dem *modus tollens* beginnt, scheint kaum vorstellbar. Die Logik ist tatsächlich eine sehr einfache und primitive Angelegenheit, man muß nur richtig damit umzugehen wissen.

Es könnte auch sein, daß es sich um eine Art Ideologie handelt, die die Logik in Mißkredit bringen will. Dafür sind allerdings die „Antinomien“, die irgendwelche widersprüchlich definierten Ausdrücke sind, nicht geeignet. Für solches müßte bei den Axiomen (also den Grundformeln und Grundregeln) der Logik angesetzt werden. Nur dort wäre es möglich, der Logik Widersprüche nachzuweisen. Gelungen ist das bisher nicht. Eben diese Ideologie ist unter Umständen aber auch dafür verantwortlich, daß es bisher keine allgemein anerkannte Lösung des Antinomienproblems gibt. Eine Lösung des Problems scheint in gewissen Kreisen einfach nicht erwünscht zu sein.

Um nicht mißverstanden zu werden:

Fachliches Mißtrauen gegen alles und jeden hat in der Wissenschaft von je her am besten gewirkt und zeichnet eine gesunde Wissenschaft aus. Alle oben gemachten Aussagen zum Antinomien-Problem beruhen auf der Widerspruchsfreiheit der Logik, ist diese nicht gegeben, ist alles hinfällig.

Fazit: Vom Standpunkt der Logik aus erweisen sich die sogenannten logischen „Antinomien“ als rein logisch erklärbare Fehler, entstanden aus konsequenter Nichtbeachtung einfachster logischer Regeln und Grundsätze.

Alternativ wäre auch vorstellbar, sich auf den Standpunkt zu stellen, daß die antinomischen Ausdrücke in jedem Fall gültige Ausdrücke des Kalküls sein müssen. Mit der klassischen Logik geht dieses nicht, wie oben zu sehen war, es müßte eine Logik erdacht werden, die wesentlich schwächer ist, so daß die Ableitung von Widersprüchen aus diesen Ausdrücken nicht mehr möglich ist. In dieser Axiomatik, sollte es denn möglich sein sie zu konstruieren, würde man aber wahrscheinlich wieder Ausdrücke finden können, die im Rahmen dieser schwächeren Axiomatik widersprüchlich sind. Dieses Spiel könnte man endlos weiterspielen, gewinnen würde man nichts.

5 Die Diagonalverfahren

Diagonalverfahren sind beliebte indirekte Beweisverfahren der Mathematik und der Theoretischen Informatik. Sie sind vergleichbar mit den hier besprochenen Antinomien. Im folgenden ist immer nur das sogenannte 2. Cantorsche Diagonalverfahren gemeint.

Gemeinsam ist Diagonalverfahren wie Antinomien der erzeugte Widerspruch. Es handelt sich um die Spezialisierung des inzwischen hinlänglich bekannten Ausdrucks:

$$\bigwedge_x (A(x, y) \Leftrightarrow \neg A(x, x))$$

Wir werden diesen Ausdruck in der einen oder anderen Form immer wieder finden.

5.1 Hermes

Hans Hermes beschreibt in [12], Seite 143 – 144, ein Anwendungsbeispiel eines Diagonalverfahrens (Die Fußnote gehört zum Zitat):

...

1. *Man kann nicht entscheiden, ob ein beliebiger (in einer Umgangssprache verfaßter) Satzsatz einen Algorithmus beschreibt, welcher zur Berechnung einer einstelligen Funktion für beliebige Argumente geeignet ist.*

Dabei wollen wir uns auf den intuitiven Standpunkt stellen. Wir schließen indirekt und nehmen an, daß das genannte Problem lösbar sei. Man kann nun alle möglichen Satzsätze der betrachteten Sprache nach ihrer Länge und bei gleicher Länge lexikographisch effektiv in eine Reihe $\varphi_0, \varphi_1, \varphi_2, \dots$ bringen. Nach unserer Annahme können wir jetzt diejenigen Satzsätze auffinden, welche keine einstellige Funktion liefern. Lassen wir diese Satzsätze fort, so erhalten wir eine Folge $\varphi'_0, \varphi'_1, \varphi'_2, \dots$, welche wir effektiv beliebig weit herstellen können. Jeder Satzsatz φ'_n beschreibt einen Algorithmus zur Berechnung einer einstelligen Funktion f_n , und jeder Satzsatz, welcher die Berechnung einer derartigen Funktion beschreibt, kommt in der Folge $\varphi'_0, \varphi'_1, \varphi'_2, \dots$ vor.

Nun definieren wir durch ein Diagonalverfahren eine Funktion f wie folgt:

$$f(n) = f_n(n) + 1.$$

f ist berechenbar. Um nun $f(n)$ zu berechnen, suche man zunächst den Satzsatz φ'_n auf. Dann berechne man mit Hilfe der in φ'_n gegebenen Anweisung $f_n(n)$. Zu dieser Zahl addiere man 1. Damit hat man $f(n)$. Wenn man die hiermit kurz angedeutete Berechnungsvorschrift vollständig aufschreibt, so erhält man einen Satzsatz, der einen Algorithmus zur Berechnung von f beschreibt. Dieser Satzsatz muß unter den $\varphi'_0, \varphi'_1, \varphi'_2, \dots$ vorkommen. Er sei identisch mit φ'_m . φ'_m beschreibt aber die Berechnung von f_m . Dies ergibt einen Widerspruch, weil $f(m) = f_m(m) + 1 \neq f_m(m)$ ¹.

...

¹Bei dem Beweis wird ein bestimmter Satzsatz (nämlich der, welcher zur Berechnung von f dient) definiert mit Bezugnahme auf die Gesamtheit aller Satzsätze. Das ist ein in der *klassischen* Mathematik durchaus geläufiges Verfahren, welches z.B. angewandt wird, wenn man eine bestimmte reelle Zahl durch einen Dedekindschen Schnitt unter Bezugnahme auf die Gesamtheit aller reellen Zahlen definiert. Dieses Verfahren wird z.B. verwendet bei dem Beweis des Zwischenwertsatzes der reellen Analysis.

Aus dem Widerspruch folgt nun, laut Hans Hermes, daß es unmöglich ist, zu entscheiden, ob ein beliebiger Satzsatz einen Algorithmus beschreibt, welcher zur Berechnung einer einstelligen Funktion für beliebige Argumente geeignet ist.

Das ist die Erweiterung, die die Diagonalverfahren von den Antinomien unterscheidet. Im Unterschied zu den Antinomien ist die Argumentation nicht mit dem Auffinden des Widerspruches beendet, sondern der Widerspruch wird, so scheint es zumindest, mit dem modus tollens gelöst.

Was hat Hermes also gemacht:

Er definiert zunächst eine Funktion f wie folgt:

$$\bigwedge_n (f(n) := f_n(n) + 1)$$

Unter der Annahme, daß alle einstelligen Funktionen abzählbar sind, muß es ein m geben, so daß gilt: $f(n) = f_m(n)$. Durch Einsetzen erhält man dann:

$$\bigwedge_n (f_m(n) := f_n(n) + 1)$$

Die Ähnlichkeit mit den bisher betrachteten Konstruktionen ist nicht zu übersehen. Die Widerspruchlichkeit dieser Definition, und zwar des ganzen Ausdruckes, läßt sich dann auch wieder mit den bereits bekannten Methoden zeigen (Spezialisierung nach m , modus tollens).

Es gibt also kein m , das die Bedingung der obigen Funktion erfüllen kann. Daher gibt es überabzählbar viele einstellige Funktionen, und daher ist auch nicht entscheidbar, ob ein beliebiger Satzsatz einen Algorithmus zur Berechnung einer einstelligen Funktion für beliebige Argumente beschreibt.

5.2 Einschub: Die „zirkuläre Definition“

Dieses erste Beispiel eine Diagonalverfahrens soll der Anlaß sein, sich wiederum mit „zirkulären Definitionen“ zu befassen, wie dieses bereits ansatzweise im Abschnitt 3.2 geschehen ist, als wir uns mit dem Begriff „Menge“ beschäftigt haben.

Wie definiert Hermes seine Funktion f ? Ganz so sorglos darf man nicht sein, denn eigentlich handelt es sich um einen Ausdruck 2. Stufe, nämlich:

$$\bigvee_f \bigwedge_n (f(n) := f_n(n) + 1)$$

Es wird also vorausgesetzt, daß es eine solche Funktion wirklich gibt, daß sie wohldefiniert ist. Aber ist sie das? Woher weiß man das?

Es wird eine einstellige Funktion definiert, die verschieden sein soll von allen einstelligen Funktionen aus einer bestimmten Menge. Wäre die Funktion selbst in der Menge, so wäre sie als verschieden von sich selbst definiert, was dem logischen Grundaxiom widerspricht, daß logische Objekte mit sich selbst identisch sein müssen. Diese Funktion f ist also nur dann wohldefiniert, nur dann logisch existent, wenn sichergestellt ist, daß es zu keiner derartigen Selbstanwendung kommen kann, sprich f also nicht in der Menge enthalten ist. Gerade das schließt Hermes aber nicht aus, ja er kann es gar nicht, will er nicht seinen Beweis zerstören.

Hermes behauptet lakonisch, daß f berechenbar ist. So allgemein stimmt das nicht, wie oben zu sehen ist. f ist nicht mit Sicherheit berechenbar, denn f ist zirkulär definiert.

Ganz allgemein ist nämlich die folgende Definition nicht erfüllbar:

$$\bigwedge_g (f := \neg g)$$

Dabei soll $\neg$ eine Operation sein, die g irgendwie verändert. Wenn man eine Funktion f definieren will, die von **allen** Funktionen verschieden ist, so gelingt das nicht, weil dann $f \neq f$ gelten müßte. So kann und darf man keine Funktion definieren.

Zirkel sind immer problematisch, aber nicht grundsätzlich falsch. Prädikate, die sich nicht verändern, sind vom zirkulären Standpunkt aus betrachtet kein Problem. Man muß sich aber natürlich fragen, was es bringt $x = id(x)$ zu definieren.

Anders ist das bei Prädikaten, die sich verändern, ein Zirkel muß dann nicht wohldefiniert sein, da die Prädikate dann unter Umständen als verschieden von sich selbst definiert sind.

D.h., Selbstanwendungen (Zirkel) können erfüllbar sein, sie müssen es aber nicht.

Wie kann man Zirkel bei Diagonalverfahren ausschließen?

1. Das durch ein Diagonalverfahren erzeugte neue Objekt kann bedingt durch seine Eigenschaften (seine Definition) keinesfalls in der vorgegebenen Menge liegen, aus der es erzeugt wurde. Dann ist das Objekt unter dem Gesichtspunkt der Selbstanwendung wohldefiniert, denn es tritt keine auf, es kann so aber auch kein Widerspruch auftreten.
2. Ist das zu definierende Objekt „augenscheinlich“ vom gleichen Typ, wie die Objekte aus denen es erzeugt werden soll, so muß explizit ausgeschlossen werden, daß das neue Objekt zu seiner eigenen Definition mit herangezogen wird. Sonst besteht die Gefahr, daß das Objekt nicht einwandfrei definiert wird, was die Benutzung in einem modus tollens verbietet.

Hermes definiert seine neue Funktion in Bezug auf die Gesamtheit aller Funktionen, und zwar veränderlich, so daß die Funktion auch als verschieden von sich selbst definiert wäre. Diese Feststellung ist ganz unabhängig von der Abzählbarkeit dieser wie auch immer gearteten Grundgesamtheit. Der im Beweis verwendete modus tollens hat dann aber zwei ungesicherte Prämissen, nämlich die Abzählbarkeit der einstelligen Funktionen UND die Wohldefiniertheit der Funktion f . Der Widerspruch führt dann dazu, daß gilt: Die einstelligen Funktionen sind nicht abzählbar ODER die Funktion f ist nicht wohldefiniert. Das reicht schon, um den Beweis zusammenbrechen zu lassen, denn der Verursacher des Widerspruches ist ersteinmal nicht eindeutig feststellbar.

Außerdem weiß man aber, daß die Funktion f eine veränderliche Selbstanwendung ist, also erfüllbar sein kann, aber nicht muß. Der Verursacher des Widerspruches ist mit ziemlicher Sicherheit gefunden. Die Abzählbarkeit der Funktionen ist es wohl nicht.

Beweise nach dem Prinzip der Diagonalverfahren haben einen eindeutig zirkulären Charakter. Es wird eine Funktion mit einer „Schwachstelle“ definiert. Diese Funktion ist nur unter ganz bestimmten Nebenbedingungen wohldefiniert. Dann wird eine Voraussetzung hergenommen, die genau diese „Schwachstelle“ trifft, d.h., die die notwendigen Nebenbedingungen für die Wohldefiniertheit der

Funktion geradezu negiert. Für den daraus folgenden Widerspruch wird dann aber nicht die zweifelhafte Funktionsdefinition, sondern die zusätzliche Voraussetzung verantwortlich gemacht.

Das eigentliche Problem liegt in der zirkulären Definition der verwendeten Funktion. Die zusätzliche Voraussetzung zeigt lediglich den in der Definition der Funktion verborgen liegenden Defekt auf, legt sozusagen den Finger in die offene Wunde. Ein Widerspruchsbeweis, der eine solche Funktionsdefinition verwendet, ist nicht in Ordnung, er macht nicht den Verursacher des Widerspruches verantwortlich.

Die zum Textauszug von Hermes gehörende Fußnote spricht Bände über die Methoden des Definierens. Die Anwendung eines Verfahrens durch den Hinweis zu begründen, daß das Verfahren **geläufig** sei, ist zwar, zumindest bei einigen Vertretern der Mathematikerzunft, eine erlaubte Vorgehensweise, aber Unarten sollte man nicht unnötig kultivieren.

Oben war zu sehen, daß es verschiedene Möglichkeiten gibt, etwas in Bezug auf Gesamtheiten zu definieren. Aber nicht immer führen diese Methoden zu wohldefinierten Ergebnissen. Vor Analogieschlüssen ist im allgemeinen zu warnen. Analogien sind selten so eindeutig, daß sich beliebige Eigenschaften übertragen lassen.

5.3 Eine weitere „Diagonalisierung“

Wir werden nun versuchen, ein Diagonalverfahren auf natürliche Zahlen anzuwenden. Natürliche Zahlen lassen sich in eindeutiger Weise binär darstellen. Wenn die natürlichen Zahlen abzählbar sind, dann kann man eine Matrix aufstellen, in deren Zeilen die natürlichen Zahlen in binärer Darstellung stehen, links beginnend mit 2^0 , 2^1 , ... und nachdem die höchste Potenz erfaßt wurde mit Nullen bis ins Unendliche aufgefüllt.

Nun kann man eine neue Zahl wie folgt erzeugen:

$$\bigwedge_n (B(n) := (M(n, n) + 1) \bmod 2)$$

Dabei ist $M(j, i)$ der Eintrag in der Binärmatrix, der die j -te Stelle der i -ten Zahl bezeichnet. Die Berechnungsvorschrift „negiert“ den Matrixeintrag, d.h. sie macht aus der 0 eine 1 und umgekehrt. Wenn die so erzeugte Zahl eine natürliche Zahl ist, und die natürlichen Zahlen abzählbar sind, dann muß diese Zahl in der Liste vorkommen, es sei die k -te. Dann gilt:

$$\bigwedge_n (M(n, k) := (M(n, n) + 1) \bmod 2)$$

Die Spezialisierung nach k ergibt: $M(k, k) <> M(k, k)$. Ein Widerspruch, also sind die natürlichen Zahlen überabzählbar. Ein recht überraschendes Ergebnis, und es stimmt auch nicht. Aber trotzdem ist man irritiert. Das Diagonalargument wurde korrekt angewendet. Oder etwa nicht? Es ist noch zu klären, ob die durch das Diagonalverfahren erzeugte Ziffernfolge wirklich eine natürliche Zahl ist.

Eine genauere Betrachtung ergibt, daß die so erzeugte Ziffernfolge ab einer gewissen Stelle nur noch aus 1en besteht, wohingegen die Binärdarstellung jeder natürlichen Zahl ab einer gewissen Stelle nur noch aus 0en besteht. Wir haben hier also den Fall 1 vor uns. Die erzeugte Ziffernfolge ist keine natürliche Zahl (sie ist irgendwie „hyper-natürlich“, eben unendlich groß), somit kommt sie nicht in der Liste vor. Der Beweis, daß die natürlichen Zahlen überabzählbar sind, funktioniert nicht, weil überhaupt kein Widerspruch auftritt. Die Definition ist nicht zirkulär.

5.4 Richard

Nun werden wir den Ansatz aus Richards Antinomie verwenden, um mit einem Diagonalverfahren zu zeigen, daß es überabzählbar viele Zahlausdrücke in der Deutschen Sprache gibt. Richard scheint wie selbstverständlich davon auszugehen, daß diese Ausdrücke abzählbar sind. Darauf beruht die „Antinomie“.

Behauptung: Es gibt überabzählbar viele Ausdrücke in der Deutschen Sprache, die Eigenschaften natürlicher Zahlen beschreiben.

Beweis: (indirekt) Angenommen es gibt abzählbar viele solcher Ausdrücke in der Deutschen Sprache. Dann lassen sich diese Ausdrücke numerieren und in einer unendlichen Folge, z.B. lexikographisch, anordnen:

$$(a) \quad A_1, A_2, A_3, \dots$$

Da die A_i durch ihre Indizes numeriert werden, kann man dieses auch so ausdrücken: Für zwei beliebig herausgegriffene Zahlen n und k muß entweder der Fall eintreten, daß n die durch A_k bezeichnete Eigenschaft besitzt oder daß n die durch A_k bezeichnete Eigenschaft nicht besitzt. Ist der erste Fall gegeben, so schreiben wir abkürzend „ $A_k(n)$ “, während wir für den zweiten Fall die Abkürzung „ $\sim A_k(n)$ “ benützen. Wir betrachten nun die Eigenschaft B , welche mittels der Formel „ $\sim A_n(n)$ “ (1) ausgedrückt wird. Sie sei folgendermaßen definiert:

$$\bigwedge_n (B(n) \leftrightarrow \sim A_n(n))$$

Dies ist offenbar eine in der deutschen Sprache definierte Eigenschaft die Eigenschaften natürlicher Zahlen beschreibt; denn diese Formel besagt ja: „ n hat nicht die Eigenschaft, welche durch A_n bezeichnet wird“ (2), und da laut Voraussetzung A_n ein Ausdruck der deutschen Sprache ist, so gilt dies auch von Satz (2) für den die Formel (1) nur eine Abkürzung darstellt. Die durch (1) bzw. (2) definierte Eigenschaft B muß somit, da die Folge (a) nach Voraussetzung alle deutschen Ausdrücke enthält, welche Zahleigenschaften definieren, mit einem dieser A_i zusammenfallen, d.h. es muß eine Zahl r geben, so daß für jede beliebige Zahl n die beiden Bedingungen $A_r(n)$ und $\sim A_n(n)$ zusammenfallen. Was für beliebiges n gilt, muß insbesondere für die spezielle Zahl r gelten. Es müßte also $A_r(r)$ dasselbe sein, wie $\sim A_r(r)$.

Das ist ein Widerspruch, es gibt kein solches r , das diese Bedingung erfüllen kann, daher haben wir einen Ausdruck in der Deutschen Sprache gefunden, der Eigenschaften natürlicher Zahlen beschreibt, der aber nicht in der Liste vorkommt, daher gibt es überabzählbar viele Ausdrücke in der Deutschen Sprache, die Eigenschaften natürlicher Zahlen beschreiben.

Auch das ist wieder ein merkwürdiges Ergebnis. Zweifelsohne ist B aber ein Ausdruck der Deutschen Sprache, der Eigenschaften natürlicher Zahlen beschreibt. Auf diese Art läßt sich also nichts machen. Das, durch das Diagonalverfahren erzeugte Objekt, ist vom gleichen Typ, wie die Objekte aus denen es erzeugt wurde. Deren Eigenschaften waren genügend allgemein gewählt.

Ist das so erzeugte Objekt aber auch in jedem Fall widerspruchsfrei, also wohldefiniert? Im Selbstanwendungsfall ist es das, wie oben zu sehen war, nicht. Wenn das Verfahren in einem Beweis verwendet wird, dann muß sichergestellt werden, daß dieser Fall nicht auftreten kann. Die Definition

des neuen Ausdruckes bezieht sich aber auf alle Ausdrücke des gleichen Typs, die Widerspruchsfreiheit der Definition ist nicht garantiert. Der Beweis der Überabzählbarkeit der Ausdrücke der Deutschen Sprache funktioniert nicht, weil ein ungesichertes Verfahren (zirkulär, veränderlich) angewendet wurde, der sich ergebende Widerspruch muß wieder nichts mit der Abzählbarkeit zu tun haben.

Wenn man in dem obigen „Beweis“ über die Überabzählbarkeit der Ausdrücke in der Deutschen Sprache, die Eigenschaften natürlicher Zahlen beschreiben, die Ausdrücke auf rein arithmetische Ausdrücke beschränkt, wie sich dieses gelegentlich in Veröffentlichungen finden läßt, so tritt Fall 1 in Kraft, der Beweis funktioniert natürlich nicht.

Und gelingt es, die Richardsche Definition rein arithmetisch zu bilden, so funktioniert der Beweis auch dann nicht, weil es sich dann wieder um eine unerlaubte Selbstanwendung handelt. Mögen die verwendeten arithmetischen Funktionen für sich durchaus einwandfrei definiert sein, so sind sie in der Zusammenstellung, wie sie durch die Richardsche Definition verlangt wird, nicht mehr in allen Fällen wohldefiniert, weil eine veränderliche Selbstanwendung auftritt.

5.5 Hopcroft / Ullman

In [15] findet sich auf Seite 9 ein Beweis, daß es mehr Funktionen als Algorithmen gibt, ziemlich analog zum Beweis bei Hermes. Es soll gezeigt werden, daß es nicht-berechenbare Funktionen gibt.

...

Unter der Annahme, daß es für jede berechenbare Funktion ein Programm bzw. einen Algorithmus zu ihrer Berechnung gibt, und daß jedes Programm bzw. jeder Algorithmus eine endliche Spezifikation besitzt, sind Programme nichts weiter, als endliche Zeichenketten über irgendeinem Alphabet. Daraus folgt, daß die Menge aller Programme abzählbar unendlich ist. Betrachten wir nun alle Funktionen, die die ganzen Zahlen auf 0 und 1 abbilden: Nehmen wir an, daß die Menge dieser Funktionen abzählbar unendlich sei, und daß daher eine bijektive Abbildung in die ganzen Zahlen existiert. Sei also f_i die Funktion, die der ganzen Zahl i entspricht. Dann kann zu der Funktion

$$f(n) = \begin{cases} 0 & \text{falls } f_n(n) = 1 \\ 1 & \text{sonst} \end{cases}$$

keine zugehörige ganze Zahl existieren, was ein Widerspruch ist. [Falls $f(n) = f_j(n)$, dann haben wir den Widerspruch $f(j) = f_j(j)$ und $f(j) \neq f_j(j)$.]

...

Die zirkulär veränderliche Definition von $f(n)$ ist wieder eine ungesicherte Prämisse, der Beweis, daß es überabzählbar viele Funktionen gibt, kann so nicht stehen bleiben.

Und noch etwas ist zu klären: Ist $f(n)$ überhaupt eine Funktion? Nun sieht $f(n)$ zwar so aus, wie Funktionen üblicherweise aufgeschrieben werden, aber das ist kein Kriterium, auch wenn einige das zu glauben scheinen.

Eine Funktion muß eine Abbildung, also eine rechtseindeutige Relation sein, d.h. für jeden Wert aus der Urbildmenge muß es einen **eindeutigen** Funktionswert aus der Bildmenge geben. Eine zirkulär veränderliche Definition, ist zunächst einmal mit sich selbst identisch, andererseits verlangt

$f(n)$ ist möglicherweise tatsächlich nicht berechenbar, aber das macht nichts, da $f(n)$ auch keine Funktion ist. $f(n)$ ist kein Beispiel für eine nicht berechenbare Funktion.

Vor diesem Hintergrund muß man nun auch das Original von Georg Cantor [7] betrachten:

$$E = (x_1, x_2, \dots, x_\nu, \dots),$$

Zu den Elementen von M gehören beispielsweise die folgenden drei:

„Ist $E_1, E_2, \dots, E_\nu, \dots$ irgendeine einfach unendliche Reihe von Elementen der Mannigfaltigkeit M , so gibt es stets ein Element E_0 von M , welches mit keinem E_ν übereinstimmt.“

$$\begin{array}{rcl} E_1 & = & (a_{1,1}, a_{1,2}, \dots, a_{1,\nu}, \dots), \\ E_2 & = & (a_{2,1}, a_{2,2}, \dots, a_{2,\nu}, \dots). \\ \dots & \dots & \dots \\ E_\mu & = & (a_{\mu,1}, a_{\mu,2}, \dots, a_{\mu,\nu}, \dots). \\ \dots & \dots & \dots \end{array}$$

Betrachten wir alsdann das Element

von M , so sieht man ohne weiteres, daß die Gleichung

25

für keinen positiven ganzzahligen Wert von μ erfüllt sein kann, da sonst für das betreffende μ und für alle ganzzahligen Werte von ν

$$b_\nu = a_{\mu,\nu},$$

also auch im besonderen

$$b_\mu = a_{\mu,\mu}$$

wäre, was durch die Definition von b_ν ausgeschlossen ist. Aus diesem Satze folgt unmittelbar, daß die Gesamtheit aller Elemente von M sich nicht in Reihenform: $E_1, E_2, \dots, E_\nu, \dots$ bringen läßt, da wir sonst vor dem Widerspruch stehen würden, daß ein Ding E_0 sowohl Element von M , wie auch nicht Element von M wäre.

...

Was kann man hierzu noch sagen, ohne sich zu wiederholen. Bedeutung findet der obige „Beweis“ dadurch, daß Cantor ihn als einen einfachen Beweis verwendet hat, um die Überabzählbarkeit der reellen Zahlen zu zeigen. Das geschieht üblicherweise so:

Wenn die reellen Zahlen abzählbar sind, dann kann man die reellen Zahlen aus dem Intervall $[0,1[$ in einer Matrix M erfassen, wobei der Vorkommateil (einschl. des Komma) abgeschnitten wird. Die Ziffernfolge, die durch das Diagonalverfahren erzeugt wird

$$\bigwedge_n (R(n) := (M(n, n) + 1) \bmod 10)$$

ist aufgrund der genügend allgemeinen Eigenschaften der reellen Zahlen augenscheinlich ebenfalls wieder eine reelle Zahl aus dem Intervall $[0,1[$, müßte also nach Voraussetzung in der Matrix stehen, und das führt wieder auf den bekannten Widerspruch. Also gibt es überabzählbar viele reelle Zahlen. Aber das verwendete Verfahren, um diese neue reelle Zahl zu definieren, ist wieder nicht in jedem Fall wohldefiniert, denn man darf eine Zahl nicht als möglicherweise verschieden von sich selbst definieren. Das ist dann keine Zahl mehr. Die Annahme der Abzählbarkeit der reellen Zahlen, führt aber nun wieder dazu, daß das Verfahren genau so angewendet wird. Der Beweis kann so nicht funktionieren. Mit einem Diagonalverfahren ist nicht zu zeigen, daß die reellen Zahlen überabzählbar sind.

5.7 Das Halteproblem

Die Theoretische Informatik behauptet, daß es nicht entscheidbar ist, ob eine beliebige Turingmaschine angesetzt auf einen beliebigen „input“ hält oder nicht. Bewiesen wird das mit Hilfe eines Diagonalverfahrens:

Dazu ein gekürzter und mit Kommentaren versehener Auszug aus [16], Seite 26–27/73:

...

(Eine Turingmaschine ist ein bestimmtes Maschinenkonzept, das mit Fähigkeiten heute üblicher Rechner übereinstimmt. Für eine genauere Definition siehe das obige Skript oder [15].

Zunächst wird in einem Satz 2 gezeigt, daß eine Sprache über irgendeinem Alphabet genau dann rekursiv aufzählbar ist, wenn sie akzeptierbar ist.

Eine Turingmaschine akzeptiert eine Sprache genau dann, wenn sie jedes Wort aus der Sprache akzeptiert.

Eine Turingmaschine akzeptiert ein Wort einer Sprache genau dann, wenn sie mit diesem Wort als input hält, und nicht etwa in eine Endlosschleife gerät.

Eine Sprache ist entscheidbar genau dann, wenn es eine Turingmaschine gibt, die z.B. mit YES terminiert, wenn ein Wort zu der Sprache gehört, und mit NO wenn es nicht dazugehört.

Erwähnt werden muß noch, daß wenn eine Sprache L entscheidbar ist, auch ihr Komplement $\bar{L}$ entscheidbar ist.

Das mag als Vorbemerkung und zur Einführung genügen.

Nun werden die Sprachen K_0 und K_1 definiert. Dabei ist $\rho(M)$ die Gödelisierung (Umwandlung in einen Zahlenkode) der Turingmaschine M und $\rho(w)$ die Gödelisierung des input w . Die Gödelisierung erlaubt es, mit Turingmaschinen über Turingmaschinen zu sprechen.)

...

$K_0 := \{ \rho(M)\rho(w); \text{ die Turingmaschine } M \text{ akzeptiert den input } w \}$

$K_1 := \{ \rho(M); \text{ die Turingmaschine } M \text{ akzeptiert den input } \rho(M) \}$

...

(Danach wird gezeigt, daß K_0 rekursiv aufzählbar (akzeptierbar) und daß, wenn die Sprache K_0 entscheidbar ist, alle rekursiv aufzählbaren Sprachen über beliebigem Alphabet entscheidbar sind. Sodann wird gezeigt, daß auch K_1 rekursiv aufzählbar ist. Damit gilt also: Wenn K_0 entscheidbar ist, dann sind auch K_1 und $\bar{K}_1$ entscheidbar. Ist $\bar{K}_1$ entscheidbar, so ist $\bar{K}_1$ auch rekursiv aufzählbar. Nun folgt der eigentliche Kernpunkt des Beweises, nämlich daß gezeigt wird, daß $\bar{K}_1$ nicht rekursiv aufzählbar sein kann.)

...

Angenommen $\bar{K}_1$ wäre rekursiv aufzählbar; dann gäbe es nach Satz 2 eine Turingmaschine $\bar{M}_1$, die $\bar{K}_1$ akzeptiert, d.h.

i) $\bar{M}_1$ angesetzt auf w hält $\Leftrightarrow w \notin K_1$.

Nun ist nach Definition von K_1

$\bar{K}_1 = \{ w ; (\text{ es gibt keine Turingmaschine } M \text{ mit } w = \rho(M)) \text{ oder } (\text{ es gibt eine Turingmaschine } M \text{ mit } w = \rho(M) \text{ und } M \text{ angesetzt auf } \rho(M) \text{ hält nicht}) \},$

$\Rightarrow$

ii) $\rho(M) \notin K_1 \Leftrightarrow M$ angesetzt auf $\rho(M)$ hält nicht.

...

(Hier wird die Sprache $\bar{K}_1$ definiert. Diese Sprache setzt sich aus zwei Teilen zusammen. Sie besteht zum einen aus den Wörtern w , deren Rück-Gödelisierung keine funktionsfähigen Turingmaschinen ergibt, und zweitens aus den Wörtern w , deren Rück-Gödelisierung zwar funktionsfähige Turingmaschinen ergibt, diese Turingmaschinen aber ange-

setzt auf ihre eigene Gödelisierung nicht halten. Unter ii) wird dann der erste Teil der Wörter weggelassen.)

...

Damit würde gelten:

$\overline{M_1}$ angesetzt auf $\rho(\overline{M_1})$ hält $\Leftrightarrow \rho(\overline{M_1}) \notin K_1 \Leftrightarrow \overline{M_1}$ angesetzt auf $\rho(\overline{M_1})$ hält nicht.

– ein Widerspruch, somit kann $\overline{K_1}$ nicht rekursiv aufzählbar sein. qed.

...

Dann ist $\overline{K_1}$ natürlich auch nicht entscheidbar. Hieraus wird dann im weiteren gefolgert, daß die Sprachen K_1 und damit auch K_0 nicht entscheidbar sind, und es daher nicht möglich ist, zu entscheiden, ob eine Turingmaschine bei einem beliebigen „input“ hält oder nicht.

Ist das ein ordentlicher Beweis?

Wie ist die Sprache $\overline{K_1}$ definiert? Ein Wort gehört zu dieser Sprache, wenn das Wort als Turingmaschine betrachtet, angesetzt auf sich selbst nicht hält, wenn es also nicht zu der Sprache gehört, die die Turingmaschine akzeptiert. Das ist eine zirkulär veränderliche Definition, denn die Zugehörigkeit zu einer Sprache, wird von der Nicht-Zugehörigkeit zu Sprachen abhängig gemacht, wobei der Selbstanwendungsfall betont nicht ausgeschlossen wird. Eine solche Sprachdefinition ist nicht mit Sicherheit widerspruchsfrei, der angewendete modus tollens hat also neben der rekursiven Aufzählbarkeit der Sprache $\overline{K_1}$ noch eine weitere ungesicherte Prämisse, der Beweis, das $\overline{K_1}$ nicht rekursiv aufzählbar ist, ist daher nicht in Ordnung.

Anders ausgedrückt und ohne direkten Bezug auf Turingmaschinen: Ein Wort x gehört zur Sprache K_y genau dann, wenn das Wort x nicht zur Sprache K_x gehört.

$$\bigwedge_x (x \in K_y \Leftrightarrow x \notin K_x)$$

Das ist ganz eindeutig eine zirkulär veränderliche Definition, die nicht erfüllbar sein muß. Es ist ganz klar erlaubt, und auch so gewollt, daß $x = y$ sein kann, man sieht den Widerspruch durch Spezialisierung sofort.

Die rekursive Aufzählbarkeit, um auf Turingmaschinen zurückzukommen, ist lediglich wieder bloß der Zeigefinger auf den bereits definitionsintern möglichen Widerspruch. Der Widerspruch könnte also wieder an der rekursiven Aufzählbarkeit ODER an der Definition von $\overline{K_1}$ liegen, der Beweis der Nicht-Aufzählbarkeit ist daher nicht in Ordnung.

Die Konstruktion der Sprache $\overline{K_1}$, analog zur Russellkonstruktion, legt außerdem nahe, daß es diese Menge, und damit auch die Sprache $\overline{K_1}$ gar nicht gibt. Kann es überhaupt widersprüchliche Sprachen geben? Sprachen sind hier als Mengen definiert, und Mengen müssen widerspruchsfrei und eindeutig sein, sonst sind es keine Mengen, sondern „inkonsistente Vielheiten“. Möglicherweise ist $\overline{K_1}$ tatsächlich nicht rekursiv aufzählbar, aber das ist ohne Belang, weil $\overline{K_1}$ keine Menge und damit auch keine Sprache ist.

Und erlaubt man auch widersprüchliche Sprachen, geht also von der Mengenbeschreibung ab, so besagt der sich ergebende Widerspruch im obigen „Beweis“ trivialerweise nichts, denn dann ist ganz deutlich, daß der modus tollens mindestens zwei ungesicherte Prämissen haben kann.

So ist die Unentscheidbarkeit des Halteproblems also nicht zu zeigen.

Ein ganz ähnlicher Beweis zu Unentscheidbarkeit des allgemeinen Halteproblems von Turingmaschinen findet sich in [15], Seite 191 – 193. Der Leser möge sich selbst überzeugen, daß auch dieser Beweis von dem in diesem Aufsatz vertretenen Standpunkt aus gesehen, nicht in Ordnung ist.

Wenn man bedenkt, wieviele Probleme der theoretischen Informatik durch Reduktion auf eben diesen obigen Beweis, im Sinne der Unentscheidbarkeit gelöst wurden (z.B. das Post'sche Korrespondenzproblem), so wird klar, das dort ein erheblicher Handlungsbedarf besteht.

5.8 Zusammenfassung

Die bekannten Antinomien beruhen auf Definitionen, denen sofort anzusehen ist, daß sie widersprüchlich sind. Es handelt sich meist um Ausdrücke der folgenden Form: $\bigwedge_x (A(x, y) \Leftrightarrow \neg A(x, x))$, also um getarnte Ausdrücke der Form $X \Leftrightarrow \neg X$. Die auftretenden Widersprüche sind selbst verursacht, finden ihre Ursache nicht in der Logik, daher ist das Antinomienproblem als erledigt zu betrachten, ohne das Änderungen an der Logik notwendig sind.

Ändern muß sich höchstens die Art und Weise einiger Anwender, mit Definitionen umzugehen.

Diagonalverfahren definieren zunächst Ausdrücke der Form $\bigwedge_x (B(x) \Leftrightarrow \neg A(x, x))$, wobei B mit A zusammenhängt, um dann unter der zusätzlichen Annahme, daß $\bigvee_y \bigwedge_x B(x) \Leftrightarrow A(x, y)$ gilt, den Widerspruch zu erzeugen. Das aber bereits die Definition von B zirkulär (da von A abhängig) ist, somit nicht erfüllbar sein muß, und daher für den Widerspruch verantwortlich sein kann, wird einfach übersehen oder nicht zur Kenntnis genommen. Der modus tollens ergibt kein eindeutiges Ergebnis, die Beweise sind so nicht in Ordnung. Übliche Beweise nach Diagonalverfahren beruhen gerade auf der Ausnutzung dieses modus tollens, es ist der zentrale Kern dieser Beweise.

Die Frage, ob es überhaupt möglich ist, eine z.B. unendliche Liste aller Objekte mit bestimmten Eigenschaften aufzustellen, erweist sich als völlig sekundär. Dieser konstruktivistische Einwand trifft nicht den Kern des Problems, und zurecht haben derartige Argumentationen gegen Diagonalverfahren nicht durchgeschlagen.

In Diagonalverfahren setzt die Auf- oder Abzählbarkeit von allen Objekten mit bestimmten Eigenschaften und von einem speziell definierten „Objekt“ im Besonderen die Zuordbarkeit dieses „Objektes“ zu den Objekten mit den bestimmten Eigenschaften voraus, also dessen Zirkularität. Um das Diagonalargument (Das „Objekt“ muß in der Liste vorkommen!) überhaupt ansetzen zu können, ist die Zirkularität des „Objektes“ unverzichtbar. Wäre das erzeugte Objekt nicht (offenbar) vom gleichen Typ wie die Objekte aus denen es definiert wurde, so könnte das Diagonalargument nicht angewendet werden.

Natürlich kann man Diagonalverfahren beliebig kompliziert darstellen, indem man mehrere Abbildungen hintereinanderschaltet, und so den Zirkel in der Definition beliebig gut versteckt. Letztlich wird aber in jedem Diagonalschluß der angebliche Widerspruch als Kern des ganzen Schlusses extrahiert, der Zirkel also deutlich hervorgehoben, und damit der Einsetzungsweg aufgezeigt. Diesen Weg braucht man nur zu verfolgen, die zirkuläre Veränderlichkeit der Definition aufzuzeigen (was

immer geht, sonst würde das Diagonalargument nicht einsetzbar sein), und kann so den Schluß gegen die zirkuläre Definition wenden.

Konsequenz: Diagonalverfahren sind als Beweismittel abzulehnen.

Im übrigen bewirken die Widersprüche in den Definitionen der verschiedenen Diagonalfunktionen, daß tatsächlich alles mit rechten Dingen zugeht. Es tritt tatsächlich im gewissen Sinne immer Fall 1 der Verhinderung von zirkulär veränderlichen Definitionen ein. Das durch die Diagonaldefinition erzeugte „Objekt“ ist immer von einem anderen Typ als die Objekte aus denen es erzeugt wurde. Und zwar entweder durch korrekte Definition, oder aber dadurch, daß das „Objekt“ widersprüchlich definiert ist. Die durch die Diagonaldefinitionen erzeugten „Objekte“ sind dann eben in Wirklichkeit keine Funktionen, keine Mengen, keine Sprachen, keine Zahlen, usw., obwohl es zunächst den Anschein hat.

6 Der Gödelsche Beweis

Nun beschäftigen wir uns zum Abschluß noch mit den berühmten Ergebnissen von Kurt Gödel. In gewisser Weise sind dessen Methoden mit der Antinomienproblematik, bzw. mit den Diagonalverfahren vergleichbar:

Kurt Gödels Artikel [10] ist in zwei Teile aufgeteilt, einen ziemlich gut verständlichen Einführungsteil und einen sehr formalen schwer verständlichen zweiten Teil.

6.1 Der Einführungsteil

Kurt Gödel schreibt auf den Seiten 174 – 176 (die Fußnoten gehören zum Textauszug):

...

Nun stellen wir einen unentscheidbaren Satz des Systems PM, d.h. einen Satz A , für den weder A noch $\text{non-}A$ beweisbar ist, folgendermaßen her:

Eine Formel aus PM mit genau einer freien Variablen, u. zw. vom Typus der natürlichen Zahlen (Klasse von Klassen) wollen wir ein Klassenzeichen nennen. Die Klassenzeichen denken wir uns irgendwie in eine Folge geordnet ¹¹, bezeichnen das n -te mit $R(n)$ und bemerken, daß sich der Begriff „Klassenzeichen“ sowie die ordnende Relation R im System PM definieren lassen. Sei α ein beliebiges Klassenzeichen; mit $[\alpha; n]$ bezeichnen wir diejenige Formel, welche aus dem Klassenzeichen α dadurch entsteht, daß man die freie Variable durch das Zeichen für die natürliche Zahl n ersetzt. Auch die Tripel-Relation $x = [y; z]$ erweist sich als innerhalb PM definierbar. Nun definieren wir eine Klasse K natürlicher Zahlen folgendermaßen:

$$n \in K \equiv \overline{Bew}[R(n); n]^* \quad (1)$$

(wobei $Bew\ x$ bedeutet: x ist eine beweisbare Formel). Da die Begriffe, welche im Definitionsvorkommen, sämtlich in PM definierbar sind, so auch der daraus zusammengesetzte

¹¹Etwa nach steigender Gliedersumme und bei gleicher Summe lexikographisch.

*Durch Überstreichen wird die Negation bezeichnet.

Begriff K , d.h. es gibt ein Klassenzeichen S^{12} so daß die Formel $[S; n]$ inhaltlich gedeutet besagt, daß die natürliche Zahl n zu K gehört. S ist als Klassenzeichen mit einem bestimmten $R(q)$ identisch, d.h. es gilt

$$S = R(q)$$

für eine bestimmte Zahl q . Wir zeigen nun, daß der Satz $[R(q); q]^{13}$ in PM unentscheidbar ist. Denn angenommen, der Satz $[R(q); q]$ wäre beweisbar, dann wäre er auch richtig, d.h. aber nach dem obigen q würde zu K gehören, d.h. nach (1) es würde $\overline{Bew}[R(q); q]$ gelten, im Widerspruch mit der Annahme. Wäre dagegen die Negation von $[R(q); q]$ beweisbar, so würde $\overline{n \in K}$, d.h. $Bew[R(q); q]$ gelten. $[R(q); q]$ wäre also zugleich mit seiner Negation beweisbar, was wiederum unmöglich ist.

Die Analogie dieses Schlusses mit der Antinomie Richard springt in die Augen; auch mit dem „Lügner“ besteht eine nahe Verwandtschaft¹⁴, denn der unentscheidbare Satz $[R(q); q]$ besagt ja, daß q zu K gehört, d.h. nach (1), daß $[R(q); q]$ nicht beweisbar ist. Wir haben also einen Satz vor uns, der seine eigene Unbeweisbarkeit behauptet¹⁵. Die eben auseinandergesetzte Beweismethode läßt sich offenbar auf jedes formale System anwenden, das erstens inhaltlich gedeutet über genügend Ausdrucksmittel verfügt, um die in der obigen Überlegung vorkommenden Begriffe (insbesondere den Begriff „beweisbare Formel“) zu definieren, und in dem zweitens jede beweisbare Formel auch inhaltlich richtig ist. Die nun folgende exakte Durchführung des obigen Beweises wird unter anderem die Aufgabe haben, die zweite der eben angeführten Voraussetzungen durch eine rein formale und weit schwächere zu ersetzen.

Aus der Bemerkung, daß $[R(q); q]$ seine eigene Unbeweisbarkeit behauptet, folgt sofort, daß $[R(q); q]$ richtig ist, denn $[R(q); q]$ ist ja unbeweisbar (weil unentscheidbar). Der im System PM unentscheidbare Satz wurde also durch metamathematische Überlegungen doch entschieden. Die genaue Analyse dieses merkwürdigen Umstandes führt zu überraschenden Resultaten, bezüglich der Widerspruchsfreiheitsbeweise formaler Systeme, die in Abschn. 4 (Satz XI) näher behandelt werden.

...

Dieser Artikel hat die Mathematik schwer erschüttert, denn er führte die Vorstellung eines vollständigen und widerspruchsfreien Axiomensystems für die Mathematik ad absurdum.

Das Kurt Gödel selbst auf die „Antinomie“ von Richard und den Lügner verweist, sollte nach den bisherigen Ergebnissen bedenklich stimmen.

¹²Es macht wieder nicht die geringsten Schwierigkeiten, die Formel S tatsächlich hinzuschreiben.

¹³Man beachte, daß „ $[R(q); q]$ “ (oder was dasselbe bedeutet „ $[S; q]$ “) bloß eine metamathematische Beschreibung des unentscheidbaren Satzes ist. Doch kann man, sobald man die Formel S ermittelt hat, natürlich auch die Zahl q bestimmen und damit den unentscheidbaren Satz selbst effektiv hinschreiben.

¹⁴Es läßt sich überhaupt jede epistemologische Antinomie zu einem derartigen Unentscheidbarkeitsbeweis verwenden.

¹⁵Ein solcher Satz hat entgegen dem Anschein nichts Zirkelhaftes an sich, denn er behauptet zunächst die Unbeweisbarkeit einer ganz bestimmten Formel (nämlich der q -ten in der lexikographischen Anordnung bei einer bestimmten Einsetzung), und erst nachträglich (gewissermaßen zufällig) stellt sich heraus, daß die Formel gerade die ist, in der er selbst ausgedrückt wurde.

Durch einige kleine Umformungen wird deutlicher, was in diesem Ausschnitt ausgesagt wird. Zusätzlich schreiben wir, um Mißverständnisse zu vermeiden, statt „ $Bew\ x$ “ die heute übliche Prädikatschreibweise „ $Bew(x)$ “. Das Zeichen „ $\equiv$ “ verwendet Gödel im Sinne von Definitionsgleichheit, es kann daher wie „ $:\Leftrightarrow$ “ behandelt werden. Folgende Definitionen tauchen auf:

$$\begin{array}{ll} \text{Definition von } K: & \bigwedge_n (n \in K \equiv \overline{Bew}([R(n); n])) \\ \text{Definition von } S: & \bigwedge_n (n \in K \equiv [S; n]) \\ \text{Definition von } q: & S = R(q) \\ \text{Außerdem gilt:} & \bigwedge_x (Bew(x) \Rightarrow x) \end{array}$$

Eingesetzt ergibt das als Definition für q :

$$\bigwedge_n ([R(q); n] \equiv \overline{Bew}([R(n); n]))$$

q muß so gewählt werden, daß die obige Definition gilt, d.h. es muß offenbar vorausgesetzt werden, daß es ein solches q tatsächlich gibt, also:

$$\bigvee_q \bigwedge_n ([R(q); n] \equiv \overline{Bew}([R(n); n]))$$

Die Spezialisierung dieser Definition nach q ergibt dann die Definition von $[R(q); q]$:

$$[R(q); q] \equiv \overline{Bew}([R(q); q])$$

Gödels Argumentation ist nun die folgende:

$$Bew([R(q); q]) \Rightarrow [R(q); q] \Rightarrow \overline{Bew}([R(q); q])$$

und:

$$Bew(\overline{[R(q); q]}) \Rightarrow \overline{[R(q); q]} \Rightarrow Bew([R(q); q])$$

Ist also $[R(q); q]$ beweisbar, so ergibt sich, daß $[R(q); q]$ nicht beweisbar ist. Ist das Negat von $[R(q); q]$ beweisbar, so ergibt sich, daß dann auch $[R(q); q]$ beweisbar ist.

Die übliche, bei den Antinomien verwendete Argumentation, einen Widerspruch in Gödels Definition aufzuzeigen, ist hier zunächst nicht anwendbar. Ausgelöst wird das durch das Prädikat $Bew(x)$, für das nach Gödel gilt: $\bigwedge_x (Bew(x) \Rightarrow x)$. Es sagt aus, daß alles, was (in dem Kalkül) beweisbar ist, (in dem Kalkül) auch wahr ist. Das ist gleichbedeutend mit der Behauptung einer gewissen Art von **Widerspruchsfreiheit** des zugrundeliegenden Axiomensystems.

Über die Schreibweise von „ $Bew(x) \Rightarrow x$ “ kann man streiten, aber Gödel verwendet es so, daher soll das hier ebenfalls geschehen. Über die andere Implikationsrichtung macht Gödel keine Angaben. Daher kann das Prädikat $Bew(x)$ auf der rechten Seite der Definition nicht weggenommen werden und links nicht eingeführt werden. Gödels Beweis bricht aus dem üblichen Antinomieschema aus. Er vermeidet geschickt den direkten Widerspruch in der Definition durch Verwendung eines Prädikates zweiter Stufe.

Die Definition ist nicht direkt auf den ersten Blick widersprüchlich. Das Ergebnis ist aber trotzdem reichlich unangenehm, widerspricht es doch irgendwie dem Beweisgedanken des inhaltlichen Schließens. Aber was Gödel da macht, daß waren nur Ableitungen, keine Beweise. Für Beweise mit dem modus ponens muß man mehr voraussetzen, nämlich, daß die Voraussetzungen wahr sind. Oder aber, für die minimale Abtrennungsbedingung, daß aus den Negaten der Voraussetzungen ebenfalls das Ergebnis ableitbar ist. So einfach kann man es sich also nicht machen. Aber bevor wir uns dieser Frage genauer zuwenden, einige wichtige Bemerkungen:

Merkwürdig sind die beständig auftauchenden Hinweise darauf, daß der von ihm (Gödel) behauptete „Satz“ in den betrachteten Axiomensystemen „definierbar“ ist. Er scheint dabei anzunehmen, daß es sich dann um wahre oder aber beweisbare Ausdrücke des Axiomensystems handelt. Das stimmt aber nicht. Es besteht im allgemeinen kein bijektiver Zusammenhang zwischen „definierbar“ und „wahr“ bzw. „beweisbar“. In fast allen Sprachen lassen sich Ausdrücke erzeugen (niederschreiben), die im Widerspruch zu den Axiomen stehen, die über die Grammatik der Sprache gelegt wurden und schließlich sind (logische) Kalküle i.w.S. auch nur Sprachen.

Völlig merkwürdig ist auch die Fußnote 15. Gödels Definition ist nämlich doch zirkelhaft, wie man an seiner eigenen Beschreibung und an der Formalisierung recht gut erkennen kann. Und wie man bei der Verwendung eines Allquantors von einem „gewissermaßen zufälligen Zusammentreffen“ sprechen kann, ist vollkommen unverständlich.

Wäre alles Wahre auch beweisbar, würde also auch $\bigwedge_x (x \Rightarrow Bew(x))$ gelten, so ließe sich Gödels Satz als widersprüchlich erweisen, denn dann könnte das Prädikat *Bew* auf beiden Seiten eingeführt oder weggeschafft werden:

$$[R(q); q] \equiv \overline{[R(q); q]}$$

oder auch:

$$Bew([R(q); q]) \equiv \overline{Bew([R(q); q])}$$

Daraus würde dann mit dem modus tollens, unter der Voraussetzung, daß $\bigwedge_x (x \Leftrightarrow Bew(x))$ gilt, folgen:

$$\neg \bigvee_q \bigwedge_n ([R(q); n] \equiv \overline{Bew([R(n); n])})$$

Damit wäre dann Gödels Definition als widersprüchlich erwiesen. Es könnte unter den obigen Voraussetzungen aus rein logischen Gründen kein solches widerspruchsfreies q geben. Damit wäre dann auch kein Beweis mit dem modus ponens möglich und auch kein Beweis mit der minimalen Abtrennungsbedingung.

Was bedeutet eigentlich der Ausdruck $\bigwedge_x (x \Rightarrow Bew(x))$, der benötigt wird, um den Widerspruch herzuleiten? Für alle Ausdrücke soll gelten, daß sie, wenn sie wahr sind, auch beweisbar sind. Das ist genau die Definition der **Vollständigkeit**.

Ein widerspruchsfreies Axiomensystem ist entweder vollständig oder unvollständig.

Ist es vollständig, so gilt: $\bigwedge_x (Bew(x) \Leftrightarrow x)$. D.h., der Wahrheitsbegriff ist mit dem Beweisbarkeitsbegriff identisch. In diesem Fall ist Gödels Aussage einfach widersprüchlich, und quasi als mit dem Lügner identisch zu bezeichnen (Der Begriff „wahr“ ist durch „beweisbar“ ersetzt.). Gödels Satz wäre dann nichts anderes als eine typische Antinomie.

Und ist das Axiomensystem unvollständig, so ist Gödels Ausdruck trivial, denn die Unvollständigkeit besagt ja gerade, daß es Ausdrücke in dem System gibt, die zwar wahr, aber nicht beweisbar sind. Ob Gödels Ausdruck in einem solchen unvollständigen System ein wahrer, aber nicht beweisbarer Ausdruck ist, muß noch dahingestellt sein. Gezeigt wird es jedenfalls nicht.

Der Beweis, daß ein Axiomensystem, wie z.B. die Prädikatenlogik 2. Stufe, nicht zugleich widerspruchsfrei und vollständig sein kann, wird an dieser Stelle unschlüssig. Gödels Beweis funktioniert nicht in vollständigen, widerspruchsfreien Axiomensystemen (und nur dort ist er überhaupt sinnvoll), denn der von ihm behauptete Ausdruck steht im direkten Widerspruch zu diesen Annahmen.

Gödels „Satz“ ist also entweder widersprüchlich oder trivial.

Andererseits behauptet Gödel aber, daß sein Ausdruck durch metamathematische Überlegungen doch als wahr erwiesen werden kann.

Gödel argumentiert so: Da der „Satz“ ja nicht beweisbar ist, ist seine Aussage, der „Satz“ also wahr. Also gibt es eine Wahrheit, die in dem Kalkül, obwohl formulierbar, nicht beweisbar ist.

Diese unbeweisbare Wahrheit „gibt“ es aber z.B. nur dann, wenn die Wahrheit mit Hilfe des modus ponens von den Voraussetzungen abgetrennt werden kann.

Gödel sitzt hier dem selben Fehler auf, der in der üblichen fehlerhaften Lügner–Argumentation auftritt, indem die Selbstbezüglichkeit nicht von Anfang an berücksichtigt wird.

Gödels metamathematischer „Beweis“ für die Wahrheit seines Ausdruckes findet unter Verwendung des Ausdruckes selbst statt, der wahr sein müßte, damit man beweisen kann, daß er wahr ist. Man sieht deutlich, daß so nichts zu beweisen ist. Der Ausdruck ist eine der Voraussetzungen im „Beweis“, er ist aber in einem widerspruchsfreien, vollständigen Axiomensystem widersprüchlich, kann also keine geeignete Voraussetzung sein. Und auch die minimale Abtrennungsbedingung ist im vollständigen und widerspruchsfreien System nicht anwendbar, wie schon in der Antinomien-diskussion bemerkt wurde.

Man lese einmal [14] unter Berücksichtigung der obigen Informationen. Man trifft mehrfach auf genau diese auch von Gödel verwendete Argumentation, die seit dem „Lügner“ über die Jahrtausende beständig tradiert wird.

Was also ist von Gödels Ergebnissen zu halten?

6.2 Der formale Teil

Der formale Teil zeigt einige Unterschiede in der Argumentation im Vergleich zum „unformalen“ ersten Teil.

Gödel definiert zunächst den Kalkül P , für den er die Existenz unentscheidbarer Sätze nachweisen will. Es ist der Logikkalkül der Principia Mathematica in Verbindung mit den Peano–Axiomen der natürlichen Zahlen.

Die Axiomatik der Principia Mathematica enthält unter anderem das sogenannte Reduzibilitätsaxiom, welches – roh gesprochen – etwa darauf hinausläuft, daß es möglich ist, Begriffe oder ihnen, wenn nicht dem Sinn, so doch dem Umfang nach gleichwertige, immer auch nicht-zirkulär zu definieren. Für die Annahme dieses Axioms lassen sich keine triftigen Gründe geltend machen, ausgenommen den, daß es nötig zu sein scheint. Mit der in diesem Text vertretenen Ansicht ist dieses Axiom nicht vereinbar. Die Frage kann nicht sein, ob eine Definition zirkulär ist oder nicht, sondern ob sie widersprüchlich ist oder nicht. Und es ist keinesfalls sicher, daß es für jeden Begriff eine widerspruchsfreie Definition geben muß.

Danach bildet er die Grundzeichen des Systems P in eineindeutiger Weise durch Primzahlen auf die natürlichen Zahlen ab. Die Art der Abbildung ist so, das auch endliche Reihen von Grundzeichen in eineindeutiger Weise abgebildet werden. Die dem Grundzeichen (bzw. der Grundzeichenreihe) a zugeordnete Zahl wird mit $\Phi(a)$ bezeichnet. Einer gegebenen Klasse oder Relation $R(a_1, a_2, \dots a_n)$ zwischen Grundzeichen oder Reihen von solchen wird diejenige Klasse (Relation) $R'(x_1, x_2, \dots x_n)$ zwischen natürlichen Zahlen zugeordnet, welche genau dann besteht, wenn es solche $a_1, a_2, \dots a_n$ gibt, so daß $\bigwedge_{i=1}^n x_i = \Phi(a_i)$ und $R(a_1, a_2, \dots a_n)$ gilt. So kann man auch die Gültigkeit von Relationen zwischen natürlichen Zahlen abbilden auf die Gültigkeit beliebiger Relationen des Kalküls, wie z.B. „Variable“, „Satz“, „beweisbar“, usw.

Unter diesen Grundzeichen oder -reihen, die beliebigen Zahlen entsprechen, werden einige seien, die den syntaktischen Regeln des Kalküls entsprechen und unter diesen wiederum einige, die in dem Kalkül auch beweisbar sind. Nehmen wir einmal an, daß es z.B. tatsächlich möglich ist, eine Übersetzung der Zeichen oder Zeichenreihen des Kalküls in Zahlen und eine solche Relation zwischen diesen Zahlen zu finden, die immer genau dann erfüllt ist, wenn die Relation der Beweisbarkeit (oder ist hier etwa Ableitbarkeit gemeint?) zwischen den Voraussetzungen und der Konklusion erfüllt ist. Das müßte aber eine sehr komplizierte Relation sein. Der Trick dabei ist, daß wenn der Kalkül es erlaubt, für jede beliebige Relation zwischen Zahlen herauszufinden, ob diese Relation besteht oder nicht, man auch für die Relationen des Kalküls, die eineindeutig auf Zahlen und die Relationen zwischen Zahlen abgebildet wurden, wüßte, ob die eigentliche Relation besteht oder nicht. D.h., der Kalkül wäre vollständig.

In einer Art Einschub führt Gödel dann die rekursiv definierten zahlentheoretischen Funktionen ein und beweist einige Sätze dazu.

...

Eine zahlentheoretische Funktion²⁵ $\Phi(x_1, x_2 \dots x_n)$ heißt rekursiv definiert aus den zahlentheoretischen Funktionen $\psi(x_1, x_2 \dots x_{n-1})$ und $\mu(x_1, x_2 \dots x_{n-1})$, wenn für alle $x_1 \dots x_n, k$ ²⁶ folgendes gilt:

$$\begin{aligned}\Phi(0, x_2 \dots x_n) &= \psi(x_2 \dots x_n) \\ \Phi(k+1, x_2 \dots x_n) &= \mu(k, \Phi(k, x_2 \dots x_n), x_2 \dots x_n).\end{aligned}\tag{2}$$

...

Rekursion ist heutzutage als bekannt vorauszusetzen, anzumerken ist, daß auch die logischen Operationen wie z.B. „und“, „oder“ und „nicht“ rekursiv sind.

²⁵D.h. ihr Definitionsbereich ist die Klasse der nicht negativen ganzen Zahlen (bzw. der n -tupel von solchen) und ihre Werte sind nicht negative ganze Zahlen.

²⁶Kleine lateinische Buchstaben (ev. mit Indizes) sind im folgenden immer Variable für nicht negative ganze Zahlen (falls nicht ausdrücklich das Gegenteil bemerkt ist).

Diesen Beweisen folgen mehrere Seiten mit rekursiven Definitionen. Dann schreibt Gödel auf den Seiten 186 - 187:

...

$$44. Bw(x) \equiv (n) \{0 < n \leq l(x) \longrightarrow Ax (n Gl x) \vee (E p, q) [0 < p, q < n \& Fl(n Gl x, p Gl x, q Gl x)]\} \& l(x) > 0$$

x ist eine *Beweisfigur* (eine endliche Folge von *Formeln*, deren jede entweder *Axiom* oder *unmittelbare Folge* aus zwei der vorhergehenden ist).

$$45. x B y \equiv Bw(x) \& [l(x)] Gl x = y \quad (3)$$

x ist ein *Beweis* für die *Formel* y .

$$46. Bew(x) \equiv (E y) y B x \quad (4)$$

x ist eine *beweisbare Formel*. [Bew(x) ist der einzige unter den Begriffen 1–46, von dem nicht behauptet werden kann, er sei rekursiv.]

Die Tatsache, die man vage so formulieren kann: Jede rekursive Relation ist innerhalb des Systems P (dieses inhaltlich gedeutet) definierbar, wird, ohne auf eine inhaltliche Deutung der Formeln aus P Bezug zu nehmen, durch folgenden Satz exakt ausgedrückt: Satz V: Zu jeder rekursiven Relation $R(x_1, x_2, \dots, x_n)$ gibt es ein n -stelliges *Relationszeichen* r (mit den *freien Variablen* ³⁸ $u_1, u_2, \dots, u_n$), so daß für alle Zahlen- n -tupel $(x_1 \dots x_n)$ gilt:

$$R(x_1 \dots x_n) \longrightarrow Bew \left[Sb \left(r \begin{array}{ccc} u_1 & \dots & u_n \\ Z(x_1) & \dots & Z(x_n) \end{array} \right) \right] \quad (3)$$

$$\overline{R}(x_1 \dots x_n) \longrightarrow Bew \left[Neg Sb \left(r \begin{array}{ccc} u_1 & \dots & u_n \\ Z(x_1) & \dots & Z(x_n) \end{array} \right) \right] \quad (4)$$

Wir begnügen uns hier damit, den Beweis dieses Satzes, da er keine prinzipiellen Schwierigkeiten bietet und ziemlich umständlich ist, in Umrissen anzudeuten³⁹. Wir beweisen den Satz für alle Relationen $R(x_1 \dots x_n)$ der Form: $x_1 = \Phi(x_2 \dots x_n)$ ⁴⁰ (wo Φ eine rekursive Funktion ist) und wenden vollständige Induktion nach der Stufe von Φ an. Für Funktionen erster Stufe (d.h. Konstante und die Funktion $x + 1$) ist der Satz trivial. Habe also Φ die m -te Stufe. Es entsteht aus Funktionen niedrigerer Stufe $\Phi_1 \dots \Phi_k$ durch die Operationen der Einsetzung oder der rekursiven Definition. Da für $\Phi_1 \dots \Phi_k$ nach induktiver Annahme bereits alles bewiesen ist, gibt es zugehörige *Relationszeichen* $r_1 \dots r_k$, so daß (3), (4) gilt. Die Definitionsprozesse, durch die Φ aus $\Phi_1 \dots \Phi_k$ entsteht

³⁸Die *Variablen* $u_1 \dots u_n$ können willkürlich vorgegeben werden. Es gibt z.B. immer ein r mit den *freien Variablen* 17, 19, 23, usw., für welches (3) und (4) gilt.

³⁹Satz V beruht natürlich darauf, daß bei einer rekursiven Relation R für jedes n -tupel von Zahlen aus den Axiomen des Systems P entscheidbar ist, ob die Relation R besteht oder nicht.

⁴⁰Daraus folgt sofort die Geltung für jede rekursive Relation, da eine solche gleichbedeutend ist mit $0 = \Phi(x_1 \dots x_n)$, wo Φ rekursiv ist.

(Einsetzung und rekursive Definition), können sämtlich im System P nachgebildet werden. Tut man dies, so erhält man aus $r_1 \dots r_k$ ein neues *Relationszeichen* r^{41} , für welches man die Geltung von (3), (4) unter Verwendung der induktiven Annahme ohne Schwierigkeit beweisen kann. Ein *Relationszeichen* r , welches auf diesem Wege einer rekursiven Relation zugeordnet ist⁴², soll rekursiv heißen.

...

Unter $Sb\left(a \begin{smallmatrix} b \\ v \end{smallmatrix}\right)$ (wo a eine Formel, v eine Variable und b ein Zeichen vom selben Typ wie v bedeutet) versteht man die Formel, welche aus a entsteht, wenn man darin v überall, wo es frei ist, durch b ersetzt.

$Neg x$ bedeutet die Negation von x .

$Z(x)$ ist das Zahlzeichen für die Zahl x .

$n Gl x$ ist das n -te Glied der der Zahl x zugeordneten Zahlenreihe.

$l(x)$ ist die Länge der x zugeordneten Zahlenreihe.

$Fl(x y z)$ bedeutet, daß x die unmittelbare Folge von y und z ist.

Hierzu gibt es nun bereits einiges zu sagen. Zunächst zur Rekursivität. Rekursivität erlaubt zwar eine schöne und einfache Konstruktionstechnik, sie schützt aber nicht vor Widersprüchlichkeit. Einfaches Beispiel: A sei rekursiv. Wenn A rekursiv ist, dann ist auch $\overline{A}$ rekursiv. Wenn A und B rekursiv sind, dann ist auch $A \wedge B$ rekursiv. $\Rightarrow A \wedge \overline{A}$ ist rekursiv. $A \wedge \overline{A}$ ist aber widersprüchlich, eine Antilogie.

D.h., daß Rekursivität nicht hinreichend ist, die Widerspruchsfreiheit und damit auch die Beweisbarkeit eines Ausdruckes zu garantieren.

Was behauptet nun Satz V? Eigentlich nichts anderes, als das der Kalkül stark genug ist, um damit für jede rekursive zahlentheoretische Relation nachweisen zu können, daß sie zwischen den angegebenen Zahlen besteht oder nicht.

Dabei ist die „vage“ Beschreibung des Satzes wirklich nur sehr vage und ob der Beweis von Satz V wirklich so einfach ist, sei noch dahingestellt. Aber gehen wir für den weiteren Verlauf davon aus, daß sich Satz V wirklich beweisen ließe.

...

Wir kommen nun ans Ziel unserer Ausführungen. Sei κ eine beliebige Klasse von *Formeln*. Wir bezeichnen mit $Flg(\kappa)$ (Folgerungsmenge von κ) die kleinste Menge von *Formeln*, die alle *Formeln* aus κ und alle *Axiome* enthält und gegen die Relation „*unmittelbare Folge*“ abgeschlossen ist. κ heißt ω -widerspruchsfrei, wenn es kein *Klassenzeichen* a gibt, so daß:

$$(n) \left[Sb \left(a \begin{smallmatrix} v \\ Z(n) \end{smallmatrix} \right) \varepsilon Flg(\kappa) \right] \& \left[Neg(v Gen a) \right] \varepsilon Flg(\kappa)$$

wobei v die *freie Variable* des *Klassenzeichens* a ist.

Jedes ω -widerspruchsfreie System ist selbstverständlich auch widerspruchsfrei. Es gilt aber, wie später gezeigt wird, nicht das Umgekehrte.

⁴¹Bei der genauen Durchführung dieses Beweises wird natürlich r nicht auf dem Umweg über die inhaltliche Deutung, sondern durch seine rein formale Beschaffenheit definiert.

⁴²Welches also, inhaltlich gedeutet, das Bestehen dieser Relation ausdrückt.

Das allgemeine Resultat über die Existenz unentscheidbarer Sätze lautet:

Satz VI: Zu jeder ω -widerspruchsfreien rekursiven Klasse κ von *Formeln* gibt es rekursive *Klassenzeichen* r , so daß weder $v \text{ Gen } r$ noch $\text{Neg } (v \text{ Gen } r)$ zu $\text{Flg}(\kappa)$ gehört (wobei v die *freie Variable* aus r ist).

Beweis: Sei κ eine beliebige rekursive ω -widerspruchsfreie Klasse von *Formeln*. Wir definieren:

$$\begin{aligned} Bw_\kappa(x) \equiv & (n) [n \leq l(x) \longrightarrow Ax (n \text{ Gl } x) \vee (n \text{ Gl } x) \varepsilon \kappa \vee \\ & (E p, q) \{0 < p, q < n \ \& \ Fl(n \text{ Gl } x, p \text{ Gl } x, q \text{ Gl } x)\}] \ \& \ l(x) > 0 \end{aligned} \quad (5)$$

(vgl. den analogen Begriff 44)

$$x B_\kappa y \equiv Bw_\kappa(x) \ \& \ [l(x)] \text{ Gl } x = y \quad (6)$$

$$\text{Bew}_\kappa(x) \equiv (E y) y B_\kappa x \quad (61)$$

(vgl. die analogen Begriffe 45, 46).

Es gilt offenbar:

$$(x) [\text{Bew}_\kappa(x) \sim x \varepsilon \text{Flg}(\kappa)] \quad (7)$$

$$(x) [\text{Bew}(x) \longrightarrow \text{Bew}_\kappa(x)] \quad (8)$$

Nun definieren wir die Relation:

$$Q(x, y) \equiv \overline{x B_\kappa \left[Sb \left(y \begin{smallmatrix} 19 \\ Z(y) \end{smallmatrix} \right) \right]}. \quad (81)$$

Da $x B_\kappa y$ [nach (6), (5)] und $Sb \left(y \begin{smallmatrix} 19 \\ Z(y) \end{smallmatrix} \right)$ (nach Def. 17, 31) rekursiv sind, so auch $Q(xy)$. Nach Satz V und (8) gibt es also ein *Relationszeichen* q (mit den *freien Variablen* 17, 19), so daß gilt:

$$\overline{x B_\kappa \left[Sb \left(y \begin{smallmatrix} 19 \\ Z(y) \end{smallmatrix} \right) \right]} \longrightarrow \text{Bew}_\kappa \left[Sb \left(q \begin{smallmatrix} 17 & 19 \\ Z(x) & Z(y) \end{smallmatrix} \right) \right] \quad (9)$$

$$x B_\kappa \left[Sb \left(y \begin{smallmatrix} 19 \\ Z(y) \end{smallmatrix} \right) \right] \longrightarrow \text{Bew}_\kappa \left[\text{Neg } Sb \left(q \begin{smallmatrix} 17 & 19 \\ Z(x) & Z(y) \end{smallmatrix} \right) \right] \quad (10)$$

Wir setzen:

$$p = 17 \text{ Gen } q \quad (11)$$

(p ist ein *Klassenzeichen* mit der *freien Variablen* 19) und

$$r = Sb \left(q \begin{smallmatrix} 19 \\ Z(p) \end{smallmatrix} \right) \quad (12)$$

(r ist ein rekursives *Klassenzeichen* mit der *freien Variablen* 17^{43} . Dann gilt:

$$Sb \left(p \begin{smallmatrix} 19 \\ Z(p) \end{smallmatrix} \right) = Sb \left([17 \text{ Gen } q] \begin{smallmatrix} 19 \\ Z(p) \end{smallmatrix} \right) = 17 \text{ Gen } Sb \left(q \begin{smallmatrix} 19 \\ Z(p) \end{smallmatrix} \right) = 17 \text{ Gen } r^{44} \quad (13)$$

[wegen (11) und (12)] ferner:

$$Sb \left(q \begin{smallmatrix} 17 & 19 \\ Z(x) & Z(p) \end{smallmatrix} \right) = Sb \left(r \begin{smallmatrix} 17 \\ Z(x) \end{smallmatrix} \right) \quad (14)$$

[nach (12)]. Setzt man nun in (9) und (10) p für y ein, so entsteht unter Berücksichtigung von (13) und (14):

$$\overline{x B_\kappa (17 \text{ Gen } r)} \longrightarrow \text{Bew}_\kappa \left[Sb \left(r \begin{smallmatrix} 17 \\ Z(x) \end{smallmatrix} \right) \right] \quad (15)$$

$$x B_\kappa (17 \text{ Gen } r) \longrightarrow \text{Bew}_\kappa \left[\text{Neg } Sb \left(r \begin{smallmatrix} 17 \\ Z(x) \end{smallmatrix} \right) \right] \quad (16)$$

Daraus ergibt sich:

1. $17 \text{ Gen } r$ ist nicht κ -*beweisbar*⁴⁵. Denn wäre dies der Fall, so gäbe es (nach 61) ein n , so daß $n B_\kappa (17 \text{ Gen } r)$. Nach (16) gälte also: $\text{Bew}_\kappa \left[\text{Neg } Sb \left(r \begin{smallmatrix} 17 \\ Z(n) \end{smallmatrix} \right) \right]$, während andererseits aus der κ -*Beweisbarkeit* von $17 \text{ Gen } r$ auch die von $Sb \left(r \begin{smallmatrix} 17 \\ Z(n) \end{smallmatrix} \right)$ folgt. κ wäre also widerspruchsvoll (umsomehr ω -widerspruchsvoll).
2. $\text{Neg } (17 \text{ Gen } r)$ ist nicht κ -*beweisbar*. Beweis: Wie eben bewiesen wurde, ist $17 \text{ Gen } r$ nicht κ -*beweisbar*, d.h. (nach 61) es gilt $(n) \bar{n} B_\kappa (17 \text{ Gen } r)$. Daraus folgt nach (15) $(n) \text{Bew}_\kappa \left[Sb \left(r \begin{smallmatrix} 17 \\ Z(n) \end{smallmatrix} \right) \right]$, was zusammen mit $\text{Bew}_\kappa [\text{Neg } (17 \text{ Gen } r)]$ gegen die ω -Widerspruchsfreiheit von κ verstoßen würde.

$17 \text{ Gen } r$ ist also aus κ unentscheidbar, womit Satz VI bewiesen ist.

...

Beim Beweise von Satz VI wurden keine anderen Eigenschaften des Systems P verwendet als die folgenden:

1. Die Klasse der Axiome und die Schlußregeln (d.h. die Relation „unmittelbare Folge“) sind rekursiv definierbar (sobald man die Grundzeichen in irgend einer Weise durch natürliche Zahlen ersetzt).
2. Jede rekursive Relation ist innerhalb des Systems P definierbar (im Sinne von Satz V).

Daher gibt es in jedem formalen System, das den Voraussetzungen 1,2 genügt und ω -widerspruchsfrei ist, unentscheidbare Sätze der Form $(x) F(x)$, wo F eine rekursiv definierte Eigenschaft natürlicher Zahlen ist, und ebenso in jeder Erweiterung eines

⁴³ r entsteht ja aus dem rekursiven *Relationszeichen* q durch Ersetzen einer *Variablen* durch die bestimmte Zahl (p).

⁴⁴Die Operationen *Gen*, *Sb* sind natürlich immer vertauschbar, falls sie sich auf verschiedene *Variable* beziehen.

⁴⁵ x ist κ -*beweisbar*, soll bedeuten: $x \in \text{Flg}(\kappa)$, was nach (7) dasselbe besagt wie: $\text{Bew}_\kappa(x)$.

solchen Systems durch eine rekursiv definierbare ω -widerspruchsfreie Klasse von Axiomen. Zu den Systemen, welche die Voraussetzungen 1,2 erfüllen, gehören, wie man leicht bestätigen kann, das Zermelo-Fraenkelsche und das v. Neumannsche Axiomensystem der Mengenlehre⁴⁷, ferner das Axiomensystem der Zahlentheorie, welches aus den Peanoschen Axiomen, der rekursiven Definitionen [nach Schema (2)] und den logischen Regeln besteht⁴⁸. Die Voraussetzung 1 erfüllt überhaupt jedes System, dessen Schlußregeln die gewöhnlichen sind und dessen Axiome (analog wie in P) durch Einsetzung aus endlich vielen Schemata entstehen.

...

Gödel nimmt an, daß die Formelklasse κ widerspruchsfrei ist. Dann definiert er aber eine neue Relation $Q(x, y)$ und fügt sie der Formelklasse κ hinzu. Er zeigt lediglich, daß Q rekursiv ist, nicht aber, daß Q mit den Formeln aus κ verträglich ist in dem Sinne, daß sie ein widerspruchsfreies Gesamtsystem bilden. Das 17. Gen r nicht κ -beweisbar ist beweist Gödel dann indirekt unter Verwendung dieser Definition Q . Er hat aber zwei Prämissen, nämlich auch die Verträglichkeit von Q und κ . Der modus tollens ist also nicht in Ordnung.

Das 17. Gen r κ -beweisbar ist beweist Gödel dann mit Hilfe seines vorherigen Ergebnisses, nämlich, daß 17. Gen r nicht κ -beweisbar ist.

Was ist nun mit Q selbst. Ein Teil von Q lautet: $Sb\left(y \frac{19}{Z(y)}\right)$, es wird also in der Formel y die freie Variable 19 überall dort, wo sie vorkommt, durch $Z(y)$, also das Zahlzeichen von y ersetzt. Dieses Zahlzeichen von y enthält natürlich ebenfalls wieder die freie Variable 19, die ebenfalls wieder durch $Z(y)$ ersetzt werden muß, usw. Das ergibt eine Rekursion ohne Rekursionsanfang, Q genügt somit nicht der oben stehenden Rekursionsdefinition. Q ist nicht korrekt definiert.

Das sind einige Ungereimtheiten, die Gödels Beweis von Satz VI zusammenbrechen lassen. Die Frage, ob die Prädikatenlogik 2. Stufe vollständig ist, soll hier nicht beantwortet werden. Es muß jedoch der Annahme widersprochen werden, Gödels „Beweis“ zeige in irgendeiner Art und Weise das Gegenteil.

Gedanken machen muß man sich nun auch über die Fortführungen des Gödelschen Gedankens durch Church, Kleene, Rosser, Tarski und viele andere.

Alle diese Beweise sind mit höchster Vorsicht zu genießen.

Im übrigen spielen gerade bei dem Gödelschen Satz in Bezug auf Ablehnung oder Akzeptanz des Satzes ideologische und weltanschauliche Vorstellungen eine große Rolle. Mit seiner geradezu esoterischen Auswirkungen liegt er voll im Trend der in letzter Zeit zu beobachtenden Esoterik-Welle. Der „Satz“ war und ist ein mit strengen Mitteln „bewiesenes“ Argument dafür, daß sich nicht alles mit strengen und genauen Mitteln erreichen läßt, Wasser auf die Mühlen derjenigen, die im wissenschaftsfeindlichen Lager schwimmen.

⁴⁷Der Beweis von Voraussetzung 1. gestaltet sich hier sogar einfacher als im Falle des Systems P , da es nur eine Art von Grundvariablen gibt (bzw. zwei bei J. v. Neumann).

⁴⁸Vgl. Problem III in D. Hilberts Vortrag: Probleme der Grundlegung der Mathematik. Math. Ann. 102.

7 Schlußbemerkung

Es ist klar, daß einige Thesen in diesem Text sehr provokativ sind. Wird doch einigen Leuten vorgeworfen, sie hätten ihre Hausaufgaben in Sachen „Grundlagen der Logik“ nicht richtig gemacht. Das fördert nicht unbedingt die Akzeptanz der Ergebnisse dieses Aufsatzes.

Außerdem widersprechen die Ergebnisse geradezu heiligen „Sätzen“ und Verfahren der Mathematik. Wer möchte schon gern z.B. auf Diagonalverfahren verzichten? Die sind so herrlich bequem.

Um es nocheinmal klarzustellen: Es wird nicht gesagt, daß das Halteproblem lösbar ist, daß die reellen Zahlen abzählbar sind, daß die Prädikatenlogik 2. Stufe vollständig und widerspruchsfrei ist, usw. Es wird nur gesagt: So kann man das nicht beweisen. Wenn es anders geht, nun gut. Wenn nicht, dann allerdings besteht gewaltiger Handlungsbedarf. Irren ist menschlich, man kann sich natürlich irren.

Jeden, der sich für die Antinomienfrage, sowie Diagonalverfahren und Gödels Satz interessiert, und den Wunsch hat, diesen Aufsatz in der Luft zu zerreißen, muß man bitten, sich zuvor Gedanken zu den folgenden Themen zu machen:

- Was ist ein Widerspruch und wie wirkt er sich im Kalkül aus?
- Welche Möglichkeiten bietet ein Kalkül, Widersprüche zu entdecken?
- Wie wirken sich Widersprüche insbesondere auf die Schlußregeln aus?
- Was ist eine Ableitung, was ein Beweis?
- Wie relevant ist das Ergebnis eines „Beweises“, der eine oder mehrere widersprüchliche Voraussetzungen enthält?
- Was heißt und was bedeutet es eigentlich, etwas zu definieren?
- Welche logischen Konsequenzen haben Definitionen?
- Was ist eine widersprüchliche Definition?
- Welche Beweislast haben Definitionen?
- Wie können sich zirkuläre Definitionen auswirken und was weiß man eigentlich über ihre Widerspruchsfreiheit?

Nachdem man diese Fragen für sich geklärt hat, sollte man sich die Ergebnisse des Aufsatzes nocheinmal ansehen.

Literatur

- [1] PETZINGER, JOHANN-MICHAEL VON: *Das Verhältnis von Begriffs- und Urteilslogik. Eine Untersuchung verschiedener Logikkalküle mit einem Exkurs über die Antinomien und den Intuitionismus*. Dissertation, Universität Tübingen, 1975.

- [2] FINSLER, PAUL: *Aufsätze zur Mengenlehre. Herausgegeben von Georg Unger.* Wissenschaftliche Buchgesellschaft, Darmstadt, 1975.
- [3] KUTSCHERA, FRANZ VON: *Die Antinomien der Logik. Semantische Untersuchungen.* Verlag Karl Alber, Freiburg / München, 1964.
- [4] CANTOR, GEORG: *Ueber unendliche, lineare Punktmannigfaltigkeiten Nr. 3.* Mathematische Annalen, 20:113–121, 1882.
- [5] CANTOR, GEORG: *Ueber unendliche, lineare Punktmannigfaltigkeiten Nr. 5.* Mathematische Annalen, 21:545 – 586, 1883.
- [6] CANTOR, GEORG: *Mitteilungen zur Lehre vom Transfiniten.* Zeitschrift für Philosophie und philosophische Kritik, 91:81–125, 1887.
- [7] CANTOR, GEORG: *Über eine elementare Frage der Mannigfaltigkeitslehre.* Jahresbericht der Deutschen Mathematiker Vereinigung, 1:75–78, 1892.
- [8] CANTOR, GEORG: *Beiträge zur Begründung der transfiniten Mengenlehre.* Mathematische Annalen, 46:481–512, 1895.
- [9] CANTOR, GEORG: *Gesammelte Abhandlungen mathematischen und philosophischen Inhalts.* Springer Verlag, Berlin, 1932. Herausgegeben von Ernst Zermelo.
- [10] GÖDEL, KURT: *Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme.* Monatshefte für Mathematik und Physik, 38:173–198, 1931.
- [11] STEGMÜLLER, WOLFGANG: *Unvollständigkeit und Unentscheidbarkeit.* Springer Verlag, Wien, 1959.
- [12] HERMES, HANS: *Aufzählbarkeit, Entscheidbarkeit, Berechenbarkeit.* Springer Verlag, Berlin, 1961.
- [13] NAGEL, ERNEST und JAMES R. NEWMAN: *Der Gödelsche Beweis.* R. Oldenbourg Verlag, München und Wien, 1979.
- [14] HOFSTADTER, DOUGLAS R.: *Gödel, Escher, Bach.* Ernst Klett Verlag, Stuttgart, 1979.
- [15] HOPCROFT, JOHN E. und JEFFREY D. ULLMAN: *Einführung in die Automatentheorie, Formale Sprachen und Komplexitätstheorie.* Addison-Wesley (Deutschland) GmbH, Bonn, 1988.
- [16] PRIESE, LUTZ: *Skript zur Vorlesung „Einführung in die Theoretische Informatik“.* Universität-Gesamthochschule Paderborn.

Außerdem bin ich den folgenden Personen (die nicht unbedingt meine Meinung zu diesem Thema teilen) zu Dank verpflichtet für unermüdliches Korrekturlesen, gute Tips und stundenlange Gespräche:

Detlef Mühle,
Dr. Johann-Michael von Petzinger,

OStR Gerhard Gillhoff,
Curd Bergmann,
Prof. Dr. Bruno Baron von Freytag Löringhoff und
Herrn Rainer Würgau