

Parakonsistente Logiken

Parakonsistente Logiken sind solche Logiken, die es zulassen, dass in einer Theorie bzw. allgemein in einer Aussagenmenge bezüglich einer Aussage A sowohl A als auch die Negation von A (nämlich $\neg A$) vorkommen, ohne dass sich dadurch alle beliebigen Aussagen herleiten lassen. Eine Menge von Aussagen ist *einfach inkonsistent*, wenn in ihr sowohl A als auch $\neg A$ vorkommen.¹ Parakonsistente Logiken sind also Logiken, in denen Widersprüche vorkommen können, ohne dass die betroffene Theorie oder Prämissenmenge trivial ist. Sie wäre trivial, wenn sich aus einem Widerspruch alles ableiten ließe. Denn eine Theorie, die alle Aussagen zugleich als wahr behauptet, behauptet gar nichts, da sie mit jeder Aussage auch deren Negation behauptet. In parakonsistenten Theorien (d.h. Theorien, deren Logik eine parakonsistente ist) werden zwar einige Aussagen zugleich mit ihrem Gegenteil behauptet, aber nicht alle. Deshalb sind diese Theorien inkonsistent und dennoch nicht-trivial (d.h. sie sind nicht *absolut inkonsistent*). Die Standard-Logik bzw. "klassische" Logik (d.h. die zweiwertige Aussagenlogik, wie sie sich bei Frege, Russell, Hilbert/Ackermann u.a. findet) lässt dies nicht zu.² Dies gilt dann auch für alle Logiken, die auf der Standard-Aussagenlogik basieren, wie die Standard-Prädikatenlogik, -Modallogik usw. In der Standard-Logik - aber auch im Intuitionismus, der die Bivalenz aufgibt - gilt das Prinzip ex

¹ Eine Theorie ist absolut inkonsistent, wenn in ihr jede Aussage wahr ist (d.h. die Theorie trivial ist) (vgl. zur Begrifflichkeit: Hunter, *Metalogic*, S.78). Davon zu unterscheiden wird außerdem sein, dass $A \wedge \neg A$ als *eine Aussage* vorkommt. Aussagenmengen können einfach inkonsistent sein, ohne auf diese *explizite* Weise inkonsistent zu sein. Aussagenmengen können auf explizite Weise einfach inkonsistent sein ohne absolut inkonsistent zu sein. Absolut inkonsistente Aussagenmengen sind auch immer explizit inkonsistent.

² Um Verwechslungen mit der traditionellen Logik (der Syllogistik) zu vermeiden, ziehe ich den Ausdruck "Standard-Logik" dem Ausdruck "klassische Logik" vor.

contradictione quod libet bzw. *ex falsum quod libet*, wobei die erste Ausdrucksweise angemessener ist: aus einem Widerspruch darf man auf Beliebiges schließen. Eine Weise dies auszudrücken ist:

$$(1) p \wedge \neg p \supset q$$

Aus dem Vorliegen von "p" und seiner Negation kann eine beliebige Aussage "q" abgeleitet werden.

In einem Kalkül des Natürlichen Schließens³ lässt sich (1) schnell beweisen:

- | | |
|----------------------------------|---|
| 1.<1> $p \wedge \neg p$ | AE (Annahmееinführung) |
| 2.<1> p | $\wedge B$ (Konjunktionsbeseitigung),1 |
| 3.<1> $p \vee q$ | $\vee E$ (Adjunktionseinführung),2 |
| 4.<1> $\neg p$ | $\wedge B$ (Konjunktionsbeseitigung),1 |
| 5.<1> q | $\vee B$ (Disjunktiver Syllogismus),3,4 |
| 6.<> $p \wedge \neg p \supset q$ | $\supset E$ (Konditionalisierung),1,5 ■ |

Aus der Annahme, dass " $p \wedge \neg p$ " vorliegt, ergibt sich mit elementaren Umformungsregeln, dass in Zeile (4) eine beliebige Aussage folgt. Fängt man statt mit " $p \wedge \neg p$ " damit an, dass sowohl "p" als auch " $\neg p$ " vorliegen, ergibt sich durch zweimaliges Konditionalisieren ein Beweis für

$$(2) p \supset (\neg p \supset q)$$

wobei sich beweisen lässt, dass (1) und (2) logisch äquivalent sind. Ist eine Theorie inkonsistent, dann ergibt sich in der Standard-Logik durch die Verwendung von (1) oder (2) und das Anwenden von Modus Ponens (Konditionalbeseitigung), dass die Theorie trivial ist.

Soll es inkonsistente und dennoch nicht-triviale Theorien geben, muss also die Logik geändert werden. Insbesondere dürfen (1) und (2) nicht mehr logische Wahrheiten sein. Wie sich das erreichen lässt, darin unterscheiden sich verschiedene Varianten parakonsistenter Logik. Innerhalb des parakonsistenten Ansatzes unterscheiden sich verschiedene Varianten

ten aber auch dannach, wie sie das Vorliegen von Widersprüchen beurteilen. Im Folgenden betrachten wir daher zwei Hauptmotivationen für den *schwachen* parakonsistenten Ansatz und den *starken* parakonsistenten Ansatz.

Der schwache parakonsistente Ansatz behauptet, dass wir eine Logik brauchen, die mit Widersprüchen umgehen kann, wobei Widersprüche indessen ein Übel sind, das wir nur vorübergehend hinnehmen wollen und müssen.

Der starke parakonsistente Ansatz behauptet darüber hinaus, dass es Widersprüche gibt, die unvermeidbar sind, und die, insofern sie sich beweisen lassen, *wahr sind*. Routley und Priest haben dafür das Kunstwort "dialethism" (bzw. "dialetheia" für den einzelnen wahren Widerspruch) eingeführt (OP:xx).⁴

Inkonsistente Theorien

Gemäß der Standard-Logik explodiert eine Theorie logisch, sobald ein einziger Widerspruch in ihr auftritt. Unter "Theorie" soll hier zunächst nicht mehr verstanden werden als eine Menge von Aussagen, die zumindest derart logisch verknüpft sind, dass alle logischen Konsequenzen, die aus irgendwelchen Aussagen in dieser Menge impliziert werden, selbst in dieser Menge sind. Dass eine Theorie "explodiert" heißt, dass mit dem Vorliegen eines einzigen Widerspruchs (der "Zündung" der Explosion) sich mit *ex contradictione quod libet* alle Aussagen in ihr ableiten lassen. So gehen wir aber nicht mit Theorien um. Von den meisten Theorien, die

³ vgl. Essler/Martinez, *Grundzüge der Logik*, Bd.I

⁴ Im Weiteren verfolge ich die Sprachregelung, dass "Antinomien" Aussagen sind, so dass sowohl A als auch $\neg A$ bewiesen werden können (bzw. $A \equiv \neg A$). "Paradoxien" sollen schon Aussagen oder Resultate sein, die widersprüchlich *scheinen* oder die im Widerspruch zu intuitiv stark verankerten Auffassungen stehen. Im Englischen wird beides oft als "Paradox" bezeichnet, wodurch jedoch der falsche Eindruck einer grossen Menge beweisbarer Widersprüche entsteht.

wir behaupten oder anwenden, wissen wir nicht, ob sie konsistent sind. Insbesondere gilt das von unseren eigenen Meinungssystemen. Oft stellen wir fest, dass wir zwei Aussagen für wahr gehalten haben, von denen wir nun feststellen, dass sie inkompatibel sind. Trotzdem haben wir zu der Zeit, als unser Meinungssystem inkonsistent (!) war, nicht Beliebiges für wahr gehalten. Auch jetzt halten wir bestimmte Aussagen für wahr und andere nicht, obwohl es sehr wahrscheinlich ist, dass zwischen Dingen, die wir für wahr halten, Inkompatibilitäten bestehen, die wir (bis jetzt) nur nicht festgestellt haben. Wir verfahren so, obwohl wir die Schlußregel *ex contradictione quod libet* kennen. Auch handelt es sich hier nicht um bloße logische Inkompetenz begleitet von mangelndem Wissen, dass ein Widerspruch vorliegt. Es gibt Fälle, wo wir von der Inkonsistenz einer Theorie oder unseres Meinungssystems, dass mehrere Theorien zugleich für wahr hält, wissen, ohne dass wir deshalb alles für wahr hielten. Diese Situation kennzeichnet auch oft die Wissenschaften. Theorien, die *eigentlich* falsch sind, d.h. Daten *widersprechen*, werden trotzdem für bestimmte Prognosen verwendet. Aus ihnen oder ihrer Verwendung wird nicht auf Beliebiges geschlossen. Beispielsweise (vgl. OP:152) legte Bohr als einer der ersten ein Atommodell vor. In Bohrs Theorie konnte ein Elektron sich um einen Kern bewegen, ohne Energie abzugeben. Nach Maxwells Gleichungen muss aber ein beschleunigtes Elektron, wie im Umlauf um einen Atomkern, Energie abgeben. Nur stehen sich hier indessen nicht einfach zwei Theorien gegenüber. Denn Bohr macht in seiner Theorie von Maxwells Gleichungen Gebrauch. Bohrs Theorie ist also in sich inkonsistent. Trotzdem wurde aus ihr nicht Beliebiges abgeleitet.

Wenn diese Beschreibung des Umgehens mit Theorien zutrifft, so versagt die Standard-Logik in der Modellierung dieses Prozesses. Irgendwie müssen wir uns als Standard-Logiker vorstellen, dass nur einige Teile der Logik zur Anwendung kommen, insbesondere nicht die Regeln, die zur Trivialisierung führen würden. Aber woher wissen wir, wie wir hier aus-

wählen sollen? - Der Ansatz der schwachen parakonsistenten Logik besteht nun darin, dass eine Logik vorgestellt wird, die genau solchen Situationen inkonsistenten Theoretisierens angemessen ist. Damit wird nicht beansprucht, dass *die* Logik (der vollendeten Wissenschaft) eine parakonsistente ist, sondern lediglich, dass die meisten Schlüsse, die in den beschriebenen Situationen gezogen werden, einer parakonsistenten Logik folgen müssen, soll es sich nicht um ein willkürliches Auswählen aus einem Angebot eigentlich zu starker Logikregeln handeln. Vertreter der starken parakonsistenten Position, dass *die* Logik eine parakonsistente Logik sein muss, können unser Umgehen mit inkonsistenten Meinungssystemen und Theorien entsprechend verstehen und als Bestätigung ihrer Position deuten (vgl. OP:4).

Dass wir Theorien verwenden, die inkonsistent sind, kann in ein transzendentes Argument für die Parakonsistenz überführt werden: Wenn wir inkonsistente Theorien verwenden, ohne dass unsere Meinungssysteme trivial werden, dann ist nach den Bedingungen der Möglichkeit hierfür zu fragen. Eine Bedingung könnte eine parakonsistente Logik sein.

In der Fähigkeit, inkonsistente Theorien Modellieren zu können, kann man auch die Voraussetzung sehen, philosophische Theorien, die inkonsistent sind oder sogar als inkonsistent verworfen wurden, zu rekonstruieren.

Antinomien und semantische Geschlossenheit

Das Vorliegen inkonsistenter Theorien und unser Umgehen damit sagt noch nichts über die Vermeidbarkeit von Widersprüchen. Widersprüche könnten ein zu vermeidendes Übel sein, mit dem sich Wissenschaftstheo-

retiker und Logiker nur notgedrungen abgeben müssen. Die starke Position der Parakonsistenz hingegen behauptet, dass einige Widersprüche unvermeidbar sind. Einige Widersprüche sind wahr! Das Hauptargument für diese These liegt im Problem der semantischen Geschlossenheit. Eine Sprache ist semantisch geschlossen, wenn sie über ihre eigene Semantik reden kann, wenn sich die Bedeutungen der Ausdrücke der Sprache in dieser Sprache angeben lassen. Der Umstand der semantischen Geschlossenheit bringt indessen die semantischen Antinomien mit sich: Wenn eine Sprache semantisch geschlossen ist, so kann sie nicht nur (i) über ihre eigenen Ausdrücke reden (das kann auch eine Sprache, die nur syntaktisch geschlossen ist)⁵, sondern kann (ii) ihren Ausdrücken auch semantische Eigenschaften zusprechen. Sobald eine Sprache über solche Namen ihrer Ausdrücke verfügt, kann sie den Namen einer offenen Formel (d.h. einer Satzfunktion mit einer Leerstelle, in die ein Name eingesetzt wird, um einen Aussagesatz zu bilden) in diese Formel einsetzen, d.h. sie kann eine Satzfunktion auf sich selbst

⁵ Das ist die Entdeckung, die sich aus Gödels Verfahren der Numerierung (Gödelisierung) von Ausdrücken ergibt. Jede Sprache, die in der Lage ist, die elementare Zahlentheorie auszudrücken, kann über ihre eigene Syntax (d.h. ihre Ausdrücke und deren syntaktische Eigenschaften) sprechen. Ein paar mögliche Unterscheidungen bei der Rede von Selbstbezüglichkeit seien kurz angeführt. Selbstreferenz kann auch verschiedene Weisen zustandekommen:

- (a) temporale Rückbezüglichkeit: etwas bezieht sich auf sich, indem sich eine Phase 2 auf eine Phase 1 des Sichdurchhaltenden bezieht (z.B. ein Gespräch);
- (b) partielle Selbstbezüglichkeit: etwas bezieht sich mit einem Teil von sich auf einen *anderen* Teil von sich (z.B. das semantische Vokabular einer Sprache auf die gegenstandsbezogene Ausdrücke *derselben Sprache*);
- (c) partiell-vermittelte Selbstbezüglichkeit: etwas bezieht sich mit einem Teil von sich auf das Ganze und damit auf das Teil, das die Beziehung herstellt (z.B. bezieht sich der Ausdruck "Sprache" als Teil auf das Ganze und *auf sich selbst*);
- (d) performative Selbstbezüglichkeit: etwas weist in seinem Geschehen mit einem Identifizierungsanspruch auf sich, der sich sofort bewährt (so ergeben sich performative Tautologien) oder widerlegt (so ergeben sich performative Selbstwidersprüche) (Beispiel: "ich"= derjenige, der dies hier spricht);
- (e) totale Selbstbezüglichkeit: etwas ist in sich so gedoppelt, dass es nur in dieser Form vorliegen kann und keine Teile hat, die nicht dieser Selbstbezüglichkeit dienen (traditionell ein Modell des Ichs).

In (c) tritt - gegenüber (a) und (b) - erstmals Referenz von etwas auf sich selbst auf, und nicht bloss ein Gesamt, dessen Teil es ist. (c) liegt in den semantischen Antinomien vor, in denen eine Aussage vermittelt einer Kennzeichnung bzw. eines Namens eine Selbstreferenz herstellt.

anwenden. Es gibt dann einen Satz, der von sich behauptet, dass ihm eine bestimmte Eigenschaft zukommt⁶. Insbesondere ergeben sich dann die Antinomien, wie die des Lügners:

(1) Satz (1) ist falsch.

Wenn Satz (1) wahr wäre, dann wäre (1) falsch, weil dann dem Satzsubjekt das Prädikat zukäme. Wenn (1) falsch wäre, dann wäre (1) wahr, da (1) gerade behauptet, falsch zu sein. Das heißt es gilt:

$\text{Wahr}(1) \equiv \text{Falsch}(1)$

bzw. da die Standard-Logik zweiwertig ist, es also jenseits des Wahren keinen anderen Wahrheitswert als das Falsche gibt:

$\text{Wahr}(1) \equiv \neg \text{Wahr}(1)$. Eine Antinomie ist eine Aussage A, bei der es sowohl für A selbst als auch für die Negation von A, $\neg A$, einen Beweis gibt. Um das am Beispiel der Lügner-Antinomie deutlich zu machen, können wir das Argument reformulieren:

Wenn Satz (1) wahr wäre, dann wäre (1) falsch, weil dann dem Satzsubjekt das Prädikat zukäme. Das ist aber ein Widerspruch (die Annahme der Wahrheit führt zur Falschheit), also (per Negationseinführung) gilt, daß (1) falsch ist (Beweis von $\neg A$). Wenn (1) falsch wäre, dann wäre (1) wahr, da (1) gerade behauptet, falsch zu sein. Auch das ist ein Widerspruch, also gilt wieder per Negationseinführung, dass (1) wahr ist (Beweis von A). Also ist (1) wahr und (1) ist falsch. ■

Auch aus " $\text{Wahr}(1) \equiv \neg \text{Wahr}(1)$ " läßt sich in der Standard-Logik schnell ein expliziter Widerspruch ableiten:

- 1.<1> $\text{Wahr}(1) \equiv \neg \text{Wahr}(1)$ AE
- 2.<1> $(\text{Wahr}(1) \supset \neg \text{Wahr}(1)) \wedge (\neg \text{Wahr}(1) \supset \text{Wahr}(1))$ Df. $\equiv$,1
- 3.<1> $\text{Wahr}(1) \supset \neg \text{Wahr}(1)$ $\wedge B$,2
- 4.<1> $\neg \text{Wahr}(1) \supset \text{Wahr}(1)$ $\wedge B$,2
- 5.<5> $\text{Wahr}(1)$ AE (aus Gründen der reductio)

⁶ Das ist der Gehalt des *Diagonalisierungs-Theorems* (vgl. Smullyan, Gödel's

6.<1,5> $\neg \text{Wahr}(1)$	$\supset B$ (Modus Ponens),3,5
7.<1> $\neg \text{Wahr}(1)$	$\neg E, \underline{5}, 6^7$
8.<8> $\neg \text{Wahr}(1)$	AE^8
9.<1,8> $\text{Wahr}(1)$	$\supset B, 4, 8$
10.<1> $\neg \neg \text{Wahr}(1)$	$\neg E, \underline{8}, 9$
11.<1> $\text{Wahr}(1)$	$\neg B, 10$ (Doppelte Negation)
12.<1> $\text{Wahr}(1) \wedge \neg \text{Wahr}(1) \wedge E, 7, 11$	■

D.h.: die Antinomien sind gemäß der Standard-Logik explizite Widersprüche.

Der Vertreter der starken Parakonsistenz (der Dialethist) behauptet nun, dass die Widersprüche unvermeidlich sind. Denn die natürliche Sprache erfüllt die Bedingungen für eine semantisch geschlossene Sprache. Da die Antinomien jedoch zugleich beweisbar sind, sind sie wahr. Es gibt wahre Widersprüche. Gezeigt werden muß dazu dreierlei:

- I. Die Widersprüche sind in einer korrekten Nicht-Standard-Logik beweisbar, wenn eine semantisch geschlossene Sprache zugrundegelegt wird.
- II. Wir müssen eine semantisch geschlossene Sprache verwenden.
- III. Es gibt kein befriedigendes Verfahren, die Widersprüche zu umgehen.

ad (I): Die Beweisbarkeit der Widersprüche bezieht sich auf das Argument unterhalb (1). In diesem Argument wird von Tarskis Konvention (T) Gebrauch gemacht:

(T) "p" ist wahr (in der Sprache L) genau dann, wenn p.

Incompleteness Theorems, S.24).

⁷ In den Zeilen (5) und (6) tritt nun ein Widerspruch auf, für den gemäß der Standard-Regel der Negationseinführung die Annahme, die in (5) gemacht wurde, verantwortlich gemacht wird; d.h. diese Annahme wird entlastet und negiert.

⁸ Jetzt wird derselbe Schritt (die reductio) noch einmal mit der Negation durchgeführt. (Da auch Annahme (8) für einen Widerspruch verantwortlich gemacht werden soll, können wir nicht einfach (7) verwenden.)

Der Anführung oder Übersetzung einer Aussage die Eigenschaft des Wahrseins zuzusprechen ist genau dann wahr, wenn die Aussage selbst, gebraucht man sie, wahr ist. Diese Konvention ist gegenüber einzelnen Wahrheitstheorien (und selbst gegenüber der Frage Bivalenz oder Mehrwertigkeit) neutral. Sie faßt einen Teil unseres intuitiven Wahrheitsbegriffes. Da dieser intuitive Wahrheitsbegriff auch von Nicht-Standard-Logiken erfaßt werden soll, läßt sich auch in diesen, wenn sie eine semantisch geschlossene Sprache zugrundelegen, die Antinomie des Lügners herleiten. Herleiten läßt sich also in jedem Fall:

(2) Wahr(1) genau dann, wenn $\neg$ Wahr(1).

Was das "genau dann, wenn" in dieser Formulierung und in der Konvention (T) besagt, hängt von der verwendeten Logik ab. In der Standard-Logik handelt es sich um die materiale Äquivalenz (" $\equiv$ "), in der Parakonsistenten Logik wird es sich um eine stärkere Verknüpfung von Aussagen handeln, da dort die materiale Implikation, als die Schuldige an der Trivialitätsdrohung (vgl. Kapitel 2), aufgegeben wird. Selbst wenn sich nicht in der oben vorgeführten Weise ein expliziter Widerspruch (in Konjunktionsform) aus (2) ableiten läßt, ist mit (2) gegeben, dass eine Aussage denselben Wahrheitswert wie ihre Negation besitzt. Verlangen wir nun von einem intuitiv akzeptablen Begriff der Negation:

(i) Wenn A falsch ist, ist die Negation von A wahr.

(ii) Wenn A wahr ist, ist die Negation von A falsch.⁹

Dann ergibt sich:

Ist "Wahr(1)" wahr, dann ist aufgrund von (2) (der Äquivalenz) auch " $\neg$ Wahr(1)" wahr. Ist "Wahr(1)" falsch, dann, dann ist aufgrund von (i) " $\neg$ Wahr(1)" wahr, also aufgrund von (2) auch "Wahr(1)" wahr! Das heißt: in jedem Fall sind "Wahr(1)" und " $\neg$ Wahr(1)" *beide wahr*. Das meint der Dialethist, wenn er von wahren Widersprüchen redet. Allerdings sind die

⁹ Diese Bedingungen ergeben noch nicht die Standard-Negation, da sie nicht als Bikonditionale formuliert sind, also zum Beispiel nicht das *Tertium Non Datur* (TND) mit sich bringen.

beiden nicht nur beide wahr, sondern *auch beide falsch*: Wenn "Wahr(1)" falsch ist, dann ist aufgrund von (2) auch " $\neg$ Wahr(1)" falsch. Wenn "Wahr(1)" wahr ist, ist aufgrund von (ii) " $\neg$ Wahr(1)" falsch, also aufgrund von (2) auch "Wahr(1) falsch. Beide sind also in jedem Fall falsch! Im Dialethismus gibt es damit Aussagen, die beide Wahrheitswerte haben: Widersprüche sind wahr, insofern sie wahr und falsch zugleich sind. Sie sind unvermeidlich, wenn wir (a) eine semantisch geschlossene Sprache verwenden und (b) in unserem Wahrheitsbegriff die Konvention (T) enthalten ist. Sie sind wahr, wenn wir an den Begriff der Negation die intuitiv einleuchtenden Forderungen (i) und (ii) stellen, die im Übrigen auch in der Standard-Logik erfüllt sind. Wenn "falsch" gleichbedeutend mit "nicht-wahr" ist, fallen die beiden Bedingungen (i) und (ii) mit der Standard-Definition der Negation zusammen, die folgende Wahrheitstafel hat:

A	$\neg A$
f	w
w	f

In der parakonsistenten Logik fallen "falsch" und "nicht-wahr" nicht zusammen. (Trotzdem können zumindest einige Systeme die intuitiven Anforderungen an die Negation befriedigen.¹⁰)

Für die Behauptung (I) des Dialethisten muss also "nur" gezeigt werden, dass wir auf die Konvention (T) nicht verzichten können.

¹⁰ Insbesondere lässt sich die Lügner-Antinomie begründen, da in vielen parakonsistenten Kalkülen der Satz vom ausgeschlossenen Dritten gilt (!), so dass die Disjunktion "wahr"/"nicht-wahr" vollständig und das Lügnerargument erlaubt ist. Satz (1) könnte also lauten: (1) Satz (1) ist nicht-wahr. Ebenfalls zu bedenken ist, dass die im zweiten Durchgang des Beweises verwendete Regel der Negationseinführung (Widerspruchsbeseitigung) vielen parakonsistenten Logiken zur Verfügung steht!) Abgelehnt wird vom parakonsistenten Logiker das Prinzip *ex contradictione quod libet*. Das heißt nicht, dass nun beliebige Widersprüche akzeptabel werden. Zugelassen werden nur die Widersprüche, die unvermeidlich sind. Alle vermeintlichen anderen Widersprüche werden vermieden, indem die widerspruchsträchtige Aussage negiert wird.

Um die Wahrheit der bewiesenen Widersprüche zu behaupten, kommt es darauf an, dass das System, in dem die Widersprüche bewiesen wurden, korrekt ist (d.h. nur logische Wahrheiten beweist). Ohne die Korrektheit kommen wir von der bloßen Beweisbarkeit der Widersprüche, trotzdem sie damit unvermeidbar auftreten, nicht zu ihrer Wahrheit. Der Dialethist benötigt also mindestens einen korrekten parakonsistenten Kalkül.

ad (II): Die natürlichen Sprachen, die wir sprechen, sind anscheinend semantisch geschlossene Sprachen. In ihnen lassen sich die semantischen Antinomien formulieren. Wer dies bestreitet, also behauptet, dass die natürlichen Sprachen *eigentlich* nicht semantisch geschlossen sind bzw. sein können, muss zeigen warum dies nicht so ist. Die Argumente gegen die semantische Geschlossenheit natürlicher Sprachen rühren nun selbst von den Antinomien her. Die Behauptung, natürliche Sprachen seien eben nicht semantisch geschlossen, bzw. die Behauptung, wir bräuchten in der Wissenschaft keine semantisch geschlossenen Sprachen, sollen zur Lösung der Antinomienproblematik beitragen. In der Begründung von (III) werden wir auf die Kritik dieser beiden Behauptungen zu sprechen kommen. Vorher möchte ich auf die besondere Relevanz dieser Frage für die Philosophie hinweisen.

Die Philosophie will nicht nur über das Sprechen in einer Sprache etwas sagen, sondern über Sprache und Kommunikation überhaupt, d.h. sie will universale Aussagen machen. Dazu bedarf es der entsprechenden Ausdrucksmöglichkeiten. Konsistent kann eine Sprache über ihre eigene Syntax reden, also beispielsweise eine allgemeine strukturelle Kennzeichnung davon geben, was ein wohlgeformter Ausdruck ist, was eine Ableitung, was eine Generalisierung usw.¹¹ Es könnte dann z.B. gesagt werden, dass es über der Sprachstufe, auf deren Index man sich syntaktisch beziehen kann, es noch andere Sprachstufen (nämlich Indices von

Aussagenmengen) gibt, die alle bestimmte - nämlich syntaktische Eigenschaften exemplifizieren. Syntaktisch kann also die Hierarchie "hinauf geredet" werden¹².

Aufgrund der semantischen Antinomien gilt dies indessen für semantische Ausdrücke und alle Ausdrücke, die auf semantische Ausdrücke verweisen, nicht.

Universale Theorien der Bedeutung, der Wahrheit, des Behauptens, des Wissens (insofern die letzteren beiden Begriffe in ihrer Analyse auf die Begriffe der Bedeutung und der Wahrheit verweisen) usw. wären dann konsistent nicht möglich! Theorien dieser Art sind es aber, für die sich die Philosophie interessiert. Wir wollen nicht nur wissen, wie wir vom Standpunkt einer nicht geklärten Metasprache aus die Bedeutung und die Wahrheit von Aussagen einer Objektsprache bestimmen können, sondern hätten gerne eine allgemeine Theorie darüber, was es beispielsweise heißt, einen Behauptungssatz zu verstehen.

Das Problem der Antinomien wird zwar oft an formalen Sprachen deutlich gemacht, doch unterscheiden sich diese formalen Sprachen in den Eigenschaften, die zu den Antinomien führen, nicht von natürlichen Sprachen. Auch dass wir die Ergebnisse und Analysen sowie die Regeln der formalen Logik auf natürliche Sprachen anwenden können und müssen, wollen wir z.B. Prozesse des Schließens verständlich machen, zeigt, dass es hier keinen relevanten Unterschied gibt.

Wir reden nicht nur de facto über Eigenschaften von allen Sprachen (so etwa selbst in empirischen Theorien wie der Transformationsgrammatik), vielmehr ist es auch zweifelhaft, ob es zwei völlig inkommensurable Sprachen geben könnte. Ein Verwenden von etwas, das vermeintlich eine völlig inkommensurable Sprache wäre, könnte von uns gar nicht als

¹¹ Dies ist möglich aufgrund z.B. der Gödelisierung der Sprache, insofern sie die Arithmetik enthält (bzw. schon das Teil-System **Q** (vgl. Boolos/Jeffrey, *Computability and Logic*, S.157ff.)).

¹² Darin liegt auch die tiefere philosophische Signifikanz von Carnaps Versuch in *Die logische Syntax der Sprache*, Philosophie allein als Syntax zu verstehen.

Sprechen identifiziert werden.¹³ *Unser Begriff* der Sprache scheint also Einheitlichkeit zuzulassen oder sogar zu verlangen. Dann müßte es aber so etwas geben wie den Inbegriff der Bestimmungen, die Sprachlichkeit kennzeichnen. Diese Bestimmungen kämen allem zu, was überhaupt als "Sprache" gekennzeichnet werden kann. Eine Erläuterung dieser Bestimmungen wäre eine universale bzw. transzendente¹⁴ Sprachtheorie. Die Regeln einer Sprache sind relativ zu dieser Sprache synthetisch a priori. Gäbe es nun einen transzendentalen Sprachrahmen, dann sind die darin vorfindlichen Bestimmungen nicht allein synthetisch a priori für eine bestimmte Sprache, sondern überhaupt synthetisch a priori.¹⁵

Aussagen solcher transzendentalen Theorien scheinen aber unmöglich zu sein: Die Hierarchie-Lösung der Antinomien verortet *jede* Aussage auf eine Sprachstufe, wobei im Falle des Auftretens semantischen Vokabulars die Stufen strikt getrennt sind. Wenn nun eine semantisch-"infizierte" Aussage vorgäbe, über alle Sprachstufen einer Sprache - oder sogar über alle Sprachen - zu reden, müßte *sie selbst* zugleich auf einer der Sprachstufen sein und über *diese* Sprachstufe *und über ihre höherstufigen Nachfolger* reden. Da dies aber in der Hierarchie-Konstruktion nicht zulässig ist, kann es solche Aussagen nicht geben. Wäre das Argument bezüglich der Existenz eines transzendentalen Begriffsrahmens, der den Inbegriff der Bestimmungen der Sprachlichkeit umfaßt, zutreffend, dann gäbe es zwar diesen Inbegriff der Bestimmungen, aber von ihm könnte nichts sinnvoll ausgesagt werden!

¹³ Vgl. Davidson, "On the Very Idea of a Conceptual Scheme"; von diesem ersten Argumentationsschritt Davidsons (Inkommensurabilität abzuweisen) muss sein zweiter - zurückzuweisender - Argumentationsschritt, dass das Fehlen einer Pluralität von Begriffsrahmen auch verbiete von *einem* Begriffsrahmen zu reden, unterschieden werden.

¹⁴ Transzendental wäre diese Theorie insofern, als sie sich mit den Bedingungen befaßt dafür, dass etwas überhaupt als Sprache identifiziert werden kann. Diese Untersuchung richtet sich auf unsere sprachliche Erkenntnisweise und nicht auf die Erkenntnis von Gegenständen selbst.

¹⁵ vgl. dazu auch Bonjour, *The Structure of Empirical Knowledge*, Anhang A.

Von diesem Ergebnis muß der Vertreter der Hierarchie-Lösung nur noch einen kleinen Schritt tun, um vorzudringen zu Wittgensteins Metapher von der Leiter, die da selbst unmöglich, wegzustoßen sei – oder zu verschiedenen Formen des Raunens über das Unsagbare.¹⁶

Man mag aber das Aussein auf eine transzendente Philosophie zurückweisen. Wenn sich Philosophie darauf beschränkt, vorliegende oder zu konstruierende Sprachen in ihrer logischen Struktur und Semantik zu rekonstruieren oder zu normieren, dann lassen sich damit viele (z.B. wissenschaftstheoretische) Anliegen der Philosophie verwirklichen, ohne in Konflikt mit der Hierarchie-Lösung zu gelangen. Es lassen sich beispielsweise allgemeine und universelle Aussagen über Sprachen und Theorien eines bestimmten Types (etwa Theorien 1.ter Stufe) machen. Bezüglich eines entsprechenden Begriffes (z.B. einer empirischen Theorie 1.ter Stufe) kann der Kommensurabilitätsanspruch verteidigt werden. Wenn die so beschreibbaren Sprachen diejenigen sind, die für die jeweiligen Anliegen der Epistemologie oder Wissenschaftstheorie ausreichen, wozu dann noch der universale oder gar transzendente Anspruch? Das Verfügen über nicht-stufige semantische Begriffe würde sich in unserem Vermögen zeigen, bei Bedarf neue Sprachen bzw. Sprachstufen einzuführen. Solange es aber keinen Bedarf gibt - und dieser kann eigentlich nur aus einem universalen theoretischen Anspruch konstant aufrechterhalten werden -, warum soll man sich um Universalität kümmern?

Erkenntnistheoretiker und Semantiker schreiben indessen auch ungeachtet der klassischen Analyse der Antinomien und deren Konsequenzen an universalen Theorien der Erkenntnis und Bedeutung. Die Bemerkung, dass sich ihre Theorie gerade nicht auf die Sprache beziehe, in der die Theorie vorgelegt wird, findet sich bei ihnen nicht. Entweder sind all diese Bemühungen in ihrem universalen Anspruch zum Scheitern verurteilt

¹⁶ Bei dieser Bemerkung handelt es sich nicht um eine bloße Polemik, sondern um den Verweis auf ein Phänomen unter Metalogikern (vgl. z.B. das letzte Kapitel in Brendels

oder mit der Hierarchie-Lösung der Antinomien stimmt etwas grundsätzlich nicht. - Schon für die Selbstvergewisserung einer transzendentalen Philosophie gilt es daher, nach Alternativen Ausschau zu halten.

ad (III): Die Motivation für die parakonsistente Logik ergibt sich auch aus den schweren Mängeln, die die bisherigen Lösungsversuche der Antinomien aufweisen. Auf die zwei Typen von klassischen Lösungsvorschlägen sei hier deshalb eingegangen¹⁷:

(IIIa) Verweise auf Mehrwertigkeit oder Wahrheitswertlücken

(IIIb) Hierarchien von Semantiken (nicht-geschlossene Sprachen) im Stile Tarskis

ad (IIIa):

Als Lösung für die semantischen Antinomien wurde vorgeschlagen, dass eine Argumentation wie die obige zeige, dass die antinomische Aussage eben keinen der beiden Wahrheitswerte "wahr" oder "falsch" habe. Die vermeintlichen Antinomien hätten einen dritten Wahrheitswert (oder verschiedene je nach Art der semantisch geschlossenen Aussage) oder wären Aussagen, denen überhaupt kein Wahrheitswert zukäme, d.h. Aussagen, die in eine Wahrheitswertlücke fallen.¹⁸ Die Aussagen (wie (1) oben), die im Beweis der Antinomie verwendet werden, verlören damit ihre Antinomien erzeugende Wirkung.

Buch *Die Wahrheit über den Lügner*).

¹⁷ Es wird hier natürlich nicht auf alle Ansätze zur Lösung der Antinomien im einzelnen eingegangen, viele geraten aber strukturell in dieselben Probleme. Bezüglich pragmatischer Lösungen sei noch angemerkt, dass Antinomien insbesondere keine performativen Selbstwidersprüche sind, da deren Kontradiktionen ausschließlich wahr sein sollen, was die Kontradiktionen von Antinomien nicht sind. Bei den semantischen Antinomien ist der Bezug auf einen Sprecher nicht wesentlich. Sie als pragmatisch nicht wohlgeformt zu betrachten führt zu einer Methode des Typs (IIIa).

¹⁸ Einen Überblick über solche Ansätze gibt: Brendel, *Die Wahrheit über den Lügner*. Dort findet sich auch die folgende Strategie, solche Vorschläge zurückzuweisen. Dieselbe Strategie wird von Priest verfolgt (IC:15-20).

Diese Lösung kann man zunächst dahingehend kritisieren, dass sie *ad hoc* angesichts der Antinomien vorgenommen wird. Der Grund, der für den nicht-zweiwertigen Charakter von (1) angegeben wird, ist gerade, dass ansonsten eine Antinomie entstünde. Wünschenswerter wäre, wenn es eine allgemeine Theorie des Versagens der Zweiwertigkeit gäbe, deren bezüglich der antinomischen Aussagen einen Anwendungsfall abgeben. Aber selbst, wenn es eine solche allgemeinere Begründung der Einführung eines dritten Wahrheitswertes oder einer Wahrheitswertlücke gibt, können wieder Antinomien auftreten. Das Auftreten von Antinomien ist kein Spezifikum einer Semantik von zwei Wahrheitswerten. Denn betrachten wir den nächst komplexeren Fall einer dreiwertigen Semantik mit den Werten "wahr", "falsch", "unbestimmt". Wenn von der vermeintlichen Antinomie (1) gesagt wird, dass sie keinen der Standardwahrheitswerte habe, also entweder, wie im hier betrachteten Fall, als "unbestimmt" zu bewerten ist oder in eine Wahrheitswertlücke falle, dann hat (1) auf jeden Fall nicht den Wahrheitswert "wahr". Wenn (1) aber nicht den Wahrheitswert "wahr" hat, dann können wir sagen - und wir tun dies ja gerade! -, dass (1) "nicht wahr" ist. "nicht wahr" fällt zwar jetzt nicht mehr zusammen mit Falschheit, da die Zweiwertigkeit fallengelassen wurde, aber von allen Aussagen, die falsch oder unbestimmt sind (bzw. in eine Wahrheitswertlücke fallen), können wir sagen, dass sie nicht wahr sind. Bezüglich einer Dreiteilung der Bewertung gilt also:

$$(D) \text{Nicht-wahr}(x) \equiv \text{Falsch}(x) \vee \text{Unbestimmt}(x)$$

wobei anstelle des "unbestimmt" auch "wertlos" oder Ähnliches stehen könnte.

Betrachten wir jetzt die Aussage:

(3) Die Aussage (3) ist nicht wahr.

Wenn (3) wahr ist, dann ist (3) nicht wahr, da dann (3) gerade die Eigenschaft zukommt, die (3) ihr zuschreibt. Ist (3) falsch, dann ist (3) nicht wahr, da sich Wahrheit und Falschheit ausschließen, also ist (3) wahr. Ist (3) unbestimmt oder wahrheitswertlos, dann ist (3) nicht wahr, da der

dritte Wert bzw. die Wahrheitswertlücke den Wahrheitswert "wahr" ausschließt, also ist (3) wahr! Wir haben also gezeigt:

$$(4) \quad \text{Wahr}(3) \supset \text{nicht-wahr}(3)$$

$$(5) \quad \text{Falsch}(3) \supset \text{Wahr}(3)$$

$$(6) \quad \text{Unbestimmt}(3) \supset \text{Wahr}(3)$$

Aus den letzten beiden Feststellungen folgt aussagenlogisch unabhängig davon, ob die Aussagenlogik zwei- oder dreiwertig ist:

$$(7) \quad \text{Falsch}(3) \vee \text{Unbestimmt}(3) \supset \text{Wahr}(3)$$

Und mit dem Prinzip (D) folgt aus Aussage (7):

$$(8) \quad \text{Nicht-Wahr}(3) \supset \text{Wahr}(3)$$

also zusammen mit unserem ersten Ergebnis (4) insgesamt:

$$(9) \quad \text{Nicht-Wahr}(3) \equiv \text{Wahr}(3)$$

eine Antinomie! Selbst wenn ein Wahrheitswert "nicht-wahr" nicht auf der Sprachebene eingeführt wird wie die ursprüngliche Trichotomie, so läßt er sich spätestens auf der nächsten metasprachlichen Reflexionsebene ausgehend von der Trichotomie definieren. Selbst wenn es durch Mehrwertigkeit also gelingt, auf einer Ebene der Semantik der Objektsprache eine Antinomie zu vermeiden, so läßt sich auf der nächst höheren Ebene wieder eine Antinomie erzeugen. Davor kann sich der Vertreter der Mehrwertigkeit nur schützen, wenn er behauptet, dass die Sprache, in der wir dieses Rasonieren durchführen semantisch stärker ist als die Sprache, in der er die Bewertungstrichotomie eingeführt hat. "Semantisch stärker" meint dabei, dass sie eine von dieser Sprachstufe getrennte Metasprache ist, so dass der Ausdruck "nicht-wahr" sich nicht auf derselben Sprachebene befindet wie "falsch" und "unbestimmt", womit Prinzip (D) und das folgende Argument blockiert wären. Damit transformiert sich der Lösungsvorschlag der Mehrwertigkeit allerdings in

den Vorschlag einer Hierarchie der Metasprachen (IIIb), auf den wir gleich zu sprechen kommen.¹⁹

Bezüglich der Lösung des Antinomienproblems sind wir also keinen Schritt weiter! Insbesondere ist zu bedenken, dass sich das gerade durchgeführte Rasonieren in der Umgangssprache durchführen lässt – wir haben es ja getan –, so dass die Mehrwertigkeit für die Umgangssprache keine Lösung des Antinomienproblems bringen kann. Und das Ausweichen auf eine Hierarchie von Metasprachen, die sich *anders* verhalten soll als das gerade umgangssprachlich vorgeführte Rasonieren zeigt außerdem, dass man *nicht mehr* über die Ausdrucksmöglichkeiten der Umgangssprache redet!²⁰

Im Übrigen gibt es Antinomien, die nicht das *Tertium Non Datur* verwenden (wie Berrys Antinomie), so dass auch eine mehrwertige bzw. nicht-bivalente Objektsprache keine Lösung bietet:

Berrys Antinomie, die den semantischen Begriff des Benennens betrifft, argumentiert wie folgt:

1. Deutsch hat ein endliches Vokabular.

also:

2. Gibt es im Deutschen nur eine endliche Anzahl von singulären Termen mit weniger als 100 Buchstaben.

aber:

3. Es gibt unendlich viele natürliche Zahlen.

also:

4. Gibt es nur eine endliche Anzahl von natürlichen Zahlen, die von einem singulären (deutschen) Term mit weniger als 100 Buchstaben benannt werden.

also:

¹⁹ Diesen Weg ist z.B. Kripkes mehrwertiger Lösungsvorschlag gegangen (vgl. Kripke, "Outline of a Theory of Truth").

²⁰ Zu einer allgemeinen Diagnose des Versagens solcher mehrwertigen Ansätze vgl. (IC:29).

5. Gibt es eine kleinste natürliche Zahl, die nicht von einem singulären (deutschen) Term mit weniger als 100 Buchstaben benannt wird.

Betrachten wir die kleinste natürliche Zahl, die nicht von einem singulären (deutschen) Term mit weniger als 100 Buchstaben benannt wird. Dann gilt per definitionem:

6. Die kleinste natürliche Zahl, die nicht von einem singulären (deutschen) Term mit weniger als 100 Buchstaben benannt wird, wird nicht von einem singulären (deutschen) Term mit weniger als 100 Buchstaben benannt.

Aber gerade das haben wir gerade getan mit dem Ausdruck "die kleinste natürliche Zahl, die von einem singulären Term mit weniger als 100 Buchstaben benannt wird". Dieser Ausdruck ist ein singulärer deutscher Term (eine Kennzeichnung). Also:

7. Die kleinste natürliche Zahl, die nicht von einem singulären (deutschen) Term mit weniger als 100 Buchstaben benannt wird, wird von einem singulären (deutschen) Term mit weniger als 100 Buchstaben benannt.

Eine Antinomie! (6) und (7) widersprechen sich. Der Satz vom ausgeschlossenen Dritten wird in dieser Argumentation nicht verwendet (vgl. IC:20,formal:32ff.).²¹

ad (IIIb):

Im letzten Abschnitt schien es so, als sei die Umgangssprache eine semantisch geschlossene Sprache - aber muß das so sein? Die etablierte Lösung der semantischen Antinomien verweist auf eine Hierarchie von Metasprachen. Selbst wenn Tarski diese Konstruktion selbst nicht auf die Umgangssprache anwenden wollte, so haben andere eben das vorgeschlagen: der Umstand, dass ansonsten Antinomien drohen, Widersprüche aber zu vermeiden sind, zeige, dass auch die Umgangssprache

²¹ vgl. auch Petersen, "Dialetheism and Paradoxes of the Berry Family".

eigentlich – dem antinomischen Anschein zum Trotze – keine geschlossene Sprache ist. Diese These soll nun kritisiert werden.²²

Wäre eine natürliche Sprache (nehmen wir das Deutsche) keine semantisch geschlossene Sprache, dann müßte man innerhalb des Deutschen verschiedene (genaugenommen mindestens abzählbar unendlich viele) Sprachstufen unterscheiden. Diese Stufen unterschieden sich bezüglich des semantischen bzw. semantisch-infizierten Vokabulars, wären also strikt zu trennen. Die semantischen Ausdrücke, die wir oberflächlich betrachtet als univok verwenden, wären in Wirklichkeit als stufenrelativ zu verstehen.

Diese Auffassung ist daher schon insofern eine normierende Rekonstruktion der natürlichen Sprache, als wir in dieser keine Hinweise auf solche Stufenindices finden. Der einzige Grund, sie anzunehmen, besteht wiederum darin, dass ansonsten die Antinomien auftreten. Neben den fehlenden sprachlichen Belege, gerät dieser Ansatz jedoch in eine Reihe von Schwierigkeiten (vgl. IC:22-28):

Nehmen wir an, dass wir eine Hierarchie von Wahrheitsprädikaten haben. Jedes von ihnen besitzt dann einen Index *i*.

Wahrheitszusprechungen wären von folgender Art:

(10) "p" ist wahr-auf-der-Stufe-n

Stufenindices können Nummern sein. Im Rahmen einer Rede über die Syntax der Sprache, zu der Numerierungen zählen, müssen wir also über sie reden können. Oben wurde dies als Möglichkeit eines syntaktischen Hinaufredens in der Hierarchie beschrieben.

Betrachten wir jetzt folgende Aussage, die uns dann zur Verfügung steht:

(11) Aussage (11) ist nicht wahr-auf-der-Stufe-von-(11).

Diese Aussage muß, wenn sie in der Hierarchie vorkommt - und das sollte sie, da wir sie mit den zur Verfügung stehenden sprachlichen Mitteln gebildet haben -, eine Stufe besitzen. Nennen wir diese Stufe

²² vgl. zum folgenden auch: Priest, "Semantic Closure". Unter (ad II) haben wir schon gesehen, dass die Behauptung der Nichtgeschlossenheit große Probleme für das

"wahr-11". Dann ist "wahr-auf-der-Stufe-von-(11)" = "wahr-11". Für jedes Wahrheitsprädikat irgendeiner Stufe gilt aber die Konvention (T). Das heißt für eine Aussage auf der Stufe "wahr-11" gilt:

(12) Wahr-11("p") genau dann, wenn p.

Angewendet auf die Aussage (11) ergibt das:

(13) Wahr-11(11) genau dann, wenn (11).

Und das heißt nach obiger Erläuterung:

(14) Wahr-11(11) genau dann, wenn nicht-wahr-11(11).

Eine Antinomie! Der Umstand, über die Indices reden zu können führt im Kontext semantischen Vokabulars also wieder zu Antinomien mittels eines - sich wieder als illegitim herausstellenden - Heraufredens der Hierarchie. Der einzige Ausweg besteht nun darin, die Voraussetzung des Arguments, dass die Aussage (11) in der Hierarchie vorkommt, aufzugeben. Damit ließe sich die Aussage (11) nicht mehr formulieren. Damit ließen sich indessen überhaupt nur Aussagen formulieren, die die Hierarchie hinunter reden - im Sinne der klassischen Hierarchie-Konstruktion. Und das hat verheerende Konsequenzen.

Wir können nämlich nun nicht mehr über die Hierarchie als Ganze reden. Das ist nicht nur ein Problem für einen Transzendentalphilosophen, wie er unter (ad II) angesprochen wurde, sondern für den Konstrukteur der Hierarchie selbst. Denn er spricht ja über die Hierarchie. Aussagen wie (15) Das Wahrheitsprädikat einer Stufe n wird auf der Stufe $n+1$ definiert. wären unausdrückbar, weil in ihnen allgemein (mittels der Variablen " n ") über Stufen gesprochen wird. Offensichtlich läßt sich (15) aber als - vermeintliche(?) - Aussage des Deutschen formulieren. Und (15) muß formuliert werden, wenn die Theorie der Sprachstufen formuliert oder gerechtfertigt werden soll. Wenn Aussage von der Art von (15) unmöglich

wären, ließe sich die Theorie der Sprachstufen gar nicht ausdrücken. Das kommt einer *reductio ad absurdum* gleich.²³

Das Problem der semantischen Hierarchie ist nicht nur, wie sie ausgesagt werden könnte im Rahmen einer semantischen Theorie. Die Hierarchie soll auch die Bedeutung des semantischen Vokabulars klären. Wenn die Hierarchie aber die Bedeutung des semantischen Vokabulars erklärt, dann müßte jemand, der das semantische Vokabular versteht, entweder die Aufbauprinzipien der Hierarchie ausdrücken können, was wie wir gesehen haben unmöglich ist, oder er müßte über die gesamte Hierarchie verfügen. Auch das letztere ist jedoch unmöglich, da es in der Hierarchie mindestens abzählbar unendlich viele Sprachstufen gibt, was alle Repräsentationsmöglichkeiten des menschlichen Geistes übersteigt - oder jedenfalls allen gegenwärtigen Theorien des Geistes und seines Verhältnisses zum Gehirn zuwiderläuft. Die semantische Hierarchie kann somit auch das Verstehen des semantischen Vokabulars nicht verständlich machen.

²³ Auch wenn Metalogiker bereit sind, dies als die besondere philosophische Tiefe ihrer Theorie, die etwas sage, was unsagbar ist, aufzufassen. Einen solchen Widerspruch zuzulassen ist besser als einen Widerspruch in der Theorie zuzulassen? In *Beyond the Limits of Thought* führt Priest nach seiner Kant-Behandlung eine "Fünfte Antinomie" ein (ebd.S.110f.), die sich analog zum Problem der semantischen Hierarchie verhält. Anders als im Fall von Kants "Antinomien" (ebd. S.94-110) handelt es sich hier um eine unausweichliche Antinomie. Sie teilt die Struktur mit den kantischen Antinomien: es liegt ein unendlicher Generator vor (Wie () ist die Ursache von()"), der einen Grenzbegriff erzeugt, von dem wir reden (wollen), obwohl auf diesen wieder der Generator (von uns) angewendet wird, so dass der Grenzbegriff wieder verschwinden müßte. Im Fall der "Fünften Antinomie" ist der Generator die Operation "Gedanke von()" bezogen auf einen Gehalt (nicht einen Akt). Ausgehend von einem Gedanken erhalten wir so eine unendliche Menge von iterierten Gedanken. Die Menge dieser Gedanken sei G. G hat kein letztes Element (genauso wie die Menge der natürlichen Zahlen), und der Generator trifft - nach Voraussetzung der Kantischen Konstruktion - *innerhalb* von G zu (analog dem Satz vom zureichenden Grunde in der "Dritten Antinomie"). Also können wir zum einen *nicht über* G denken (G selbst ist transzendent), aber zum anderen tun wir dies gerade (G ist immanent): ein Widerspruch! G verhält sich genauso wie das Reden über Semantik in semantischen Hierarchien und ist jenseits der Grenzen des Denkens - eben *Beyond the Limits of Thought*. Problematisch ist natürlich auch hier die Hierarchie-Konstruktion, in welcher der Generator nur "nach unten" arbeitet. (Priest bewertet die Lage nicht eindeutig.)

Die Hierarchie-Konzeption führt ebenfalls zu Schwierigkeiten bei der Einführung einiger logischer Grundbegriffe. Denn das Konzept der Hierarchie wird nicht nur in der Semantik verwendet, sondern auch in der Mengenlehre (vgl. IC:35-47). Die kumulative Hierarchie in der Mengenlehre soll die Inkonsistenz der naiven Mengenlehre, die sich durch das unbeschränkte Aussonderungsaxiom ergibt, vermeiden. Die kumulative Hierarchie entsteht dadurch, dass ausgehend von einer möglicherweise leeren Grundmenge von Objekten durch die Operation der Potenzmengenbildung (der Bildung der Menge aller Teilmengen einer zugrundegelegten Menge) neue mächtigere (mehr Elemente enthaltende) Mengen gebildet werden. Ist beispielsweise die Grundmenge leer, $\emptyset$, so ist die Potenzmenge dieser Menge (d.h. der leeren Menge) $\{\emptyset\}$, da die leere Menge Teilmenge jeder Menge ist - es gibt in ihr ja kein Element, das in einer beliebigen anderen Menge fehlt. $\{\emptyset\}$ hat also ein Element, ist also mächtiger als die Ausgangsmenge $\emptyset$. Die Potenzmenge von $\{\emptyset\}$ ist $\{\emptyset, \{\emptyset\}\}$, da die leere Menge Teilmenge jeder Menge ist und die Einermengen der Elemente einer Menge Elemente der Potenzmenge dieser Menge sind (hier die Einermenge von $\emptyset$, die Element von unserer ersten Potenzmenge, $\{\emptyset\}$, ist). Ausgehend von der leeren Menge lassen sich so beliebig viele (mindestens überabzählbar unendlich viele²⁴) Mengen einführen.²⁵

Die Lösung der Antinomien besteht darin, dass sich Aussonderung immer auf eine Menge in dieser Hierarchie beziehen muß, beispielsweise durch das beschränkte Aussonderungsaxiom der Zermelo-Fraenkel-Axiomatisierung der Mengenlehre:

²⁴ Genaugenommen gibt es keine Grenzen der Kardinalität, vielmehr kann es sogar von dieser Konstruktionsmethode unerreichbare Kardinalzahlen geben. Die Details sind hier nicht wichtig.

²⁵ Die Hierarchie ist "kumulativ", da eine Menge einer höheren Stufe Mengen beliebiger niedrigerer Stufen enthalten kann, während in einer strikten Typentheorie nur Elemente der nächst niedrigeren Stufe zulässig sind.

(AS) Zu jeder Menge a und jeder Bedingung $F(x)$ gibt es eine Menge b , deren Elemente genau jene aus a sind, für die $F(x)$ gilt.

bzw. als Axiomenschema:

$$(AS') (\forall X)(\forall Y)(\forall z)(z \in Y \equiv z \in X \wedge P(z)).^{26}$$

Nun stellen sich wieder eine Reihe von Fragen: Gibt es Aussonderungen, die wahr sind, aber nicht in der Hierarchie vorkommen? Und entscheidender: Wie läßt sich die kumulative Hierarchie als Konzeption einführen?

Das hier auftretende Dilemma des Unsagbaren entspricht dem der semantischen Hierarchie: In der Konstruktion der Hierarchie verwenden wir den Begriff "Menge", beispielsweise wenn behauptet wird, dass wir von einer beliebigen Menge zu einer Menge höherer Kardinalität (zu einer mächtigeren Menge) durch die Potenzmengenbildung fortschreiten können. Was entspricht aber in der Hierarchie dem Begriff "Menge"? Nichts! Nichts kann diesem Begriff entsprechen: die Extension des so allgemein verwendeten Begriffs Menge müßte die Menge aller Mengen sein, aber eine Allmenge kann es in der Hierarchie nicht geben. Dies läßt sich ausgehend von der obigen Fassung des Aussonderungsaxioms schnell beweisen:²⁷

1. a sei eine beliebige Menge.

Für die auszusondernde Teilmenge verlangen wir:

$$2. b = \{x \in a \mid x \notin x\}$$

Wir wollen also die Menge der x aussondern, die nicht Selbstelemente sind.

Dann gilt nach (AS) und (2) für beliebige y :

$$3. y \in b \equiv y \in a \wedge y \notin y$$

Kann nun $b \in a$ wahr sein? Betrachten wir die beiden Fälle $b \in b$ und

²⁶ Dabei sind "X" und "Y" Variablen für Mengen, "z" mag sich auf Grundelemente oder Mengen beziehen. Zur ZF-Mengenlehre vgl. Halmos, *Naive Mengenlehre* [der Titel dieses Buches ist ein entscheidender Fehlgriff, da die ZF-Mengenlehre gerade keine naive Mengenlehre ist; Halmos will sie wohl von Typentheorien unterscheiden].

²⁷ vgl. Halmos, *Naive Mengenlehre*, S.18

$b \notin b$ (eine vollständige Disjunktion):

4. sei $b \in a$

dann ergibt sich mit $b \in b$, was die linke Seite von (3) wahr macht:

5. $b \in b \equiv b \in a \wedge b \notin b$

ein Widerspruch! Analog ergibt sich aus (4) und $b \notin b$, die zusammen die rechte Seite von (3) wahr machen:

6. $b \in b \equiv b \in a \wedge b \notin b$

derselbe Widerspruch! D.h. das die Annahme (4) falsch sein muss:

7. $b \notin a$

Nun war a aber beliebig gewählt, d.h. zu jeder Menge a gibt es eine andere Menge, die nicht Element dieser Menge a ist, also gibt es keine Allmenge. ■

Wenn es keine Allmenge gibt, dann gibt es indessen keine Extension für den Begriff "Menge (überhaupt)", der damit in der Konstruktion der kumulativen Hierarchie nicht verwendet werden könnte, wenn man überhaupt noch meint, es handele sich dann um einen sinnvollen Begriff.

Mit dem allgemeinen Begriff der Menge können auch einige andere Begriffe nicht mehr verwendet werden: das "absolute Komplement" einer Menge (unabhängig vom Rang) oder "Menge mit Elementen eines beliebig hohen Rangs". Teilbereiche der Mathematik wie die Kategorientheorie wollen aber über solche Mengen reden (vgl. IC:41f.).

Entsprechendes gilt für die Unterscheidung von Mengen und "echten Klassen", die nicht selbst Elemente sind, während Mengen auch selbst Element sein können. Zum einen wären echte Klassen außerhalb der Hierarchie anzusiedeln, wobei der einzige Grund hierfür in der Vermeidung der Antinomien liegt. Zum anderen ist nicht einzusehen, warum echte Klassen (als Entitäten mit Elementen) *überhaupt nirgendwo* Element sein *können*. Das ist eine übergeneralisierende Forderung, die sich auch wieder nur aus der Antinomienproblematik ergibt. Hier könnte

man bestenfalls eine *zweite* kumulative Hierarchie der echten Klassen anschließen - usw.

Unabhängig von diesen innermathematischen Schwierigkeiten des letzten Absatzes ergibt sich ein philosophisch entscheidender Einwand aus der Funktion der Mengenlehre für die Logik. Bei der Semantik logischer Systeme werden die Interpretationen mengentheoretisch gekennzeichnet: sie sind *Tupel* (d.h. Mengen) derart, dass die Interpretationsfunktion ausgehend von einer Grundmenge den Designatoren Extensionen oder wieder andere *Funktionen* (z.B. von möglichen Welten in *Mengen*) zuweist. Die übliche Definition der Gültigkeit von Folgerungen besagt dann:

(G) $\Sigma \models A := A$ folgt aus der Menge von Aussagen Σ genau dann, wenn jede Interpretation, die alle Elemente von Σ wahr macht, auch A wahr macht.

Nun ist die Grundmenge einer Interpretation beliebig. Sie kann also eine Menge beliebig hohen Rangs sein. Also gibt es Interpretationen beliebig hohen Rangs. Damit redet die Definition der Gültigkeit generalisierend über das ganze Universum der kumulativen Hierarchie (und über *jede* Interpretation), wobei dieses Universum aber gerade nach den Ansprüchen der Hierarchie-Konstruktion keine Menge ist. Entweder ist die Definition (G) der Gültigkeit also gemäß der kumulativen Hierarchie nicht möglich, wie aber dann sollte Gültigkeit definiert werden, da man ja nicht einfach bei der Beurteilung logischer Wahrheit einfach einige Interpretationen weglassen kann? Oder die mit (G), das wir ja anscheinend gebrauchen und verstehen, vorausgesetzte Mengenlehre ist nicht die der kumulativen Hierarchie.

Auch die kumulative Hierarchie untergräbt also die logische und sprachliche Basis, auf der sie ruhen soll. Deshalb bedarf auch die Mengenlehre einer nicht-hierarchischen Konstruktion. Da damit Antinomien auftreten werden, muss sie parakonsistent sein, soll sie nicht trivial werden.

Fazit: Der starke parakonsistente Ansatz, die Auffassung, dass es unvermeidliche und wahre Widersprüche gibt, ist m.E. ausgezeichnet motiviert durch das philosophische Interesse an semantisch geschlossenen Sprachen und dem philosophischen Ungenügen der bisherigen Lösungsvorschläge zum Antinomienproblem. Ich werde deshalb im Folgenden – zumindest zunächst – den Standpunkt des Dialethismus einnehmen und nach der Ausführung einer entsprechenden Logik fragen.

In Erinnerung zu behalten ist dabei die "Nicht-ad-hoc-Forderung": Einzelne Vorschläge müssen sich dadurch rechtfertigen, nicht nur eine Aufgabe zu lösen (z.B. eine Antinomie zu vermeiden), wie dies einigen klassischen Lösungen vom Dialethisten vorgeworfen wurde. Je weiter anwendbar sich parakonsistente Logiken anwenden lassen, um so akzeptabler wird ihre Lösung der Antinomienproblematik sein, und umgekehrt.

Zum einen gehe ich davon aus, dass der starke parakonsistente Ansatz ein berechtigtes philosophisches Anliegen vorbringt, das es zu verfolgen gilt. Ansätze und Überlegungen zur Parakonsistenz, die sich hauptsächlich formalistisch motivieren, in dem Sinne, dass Systeme entwickelt und untersucht werden aus dem Interesse heraus, was denn formal sich durchführen lässt, jenseits von einer Ankoppelung an philosophische Adäquatheits- oder Relevanzkriterien, verfolge ich deshalb nicht, vielmehr werde ich gelegentlich darauf hinweisen, warum ich sie für philosophisch eher irrelevant halte.

Zum anderen orientiere ich mich aufgrund der Kritik an anderen formalen parakonsistenten Ansätzen am relevanzlogischen Ansatz von Priest und Routley. Hier käme es darauf an, einen adäquaten Kalkül der parakonsistenten Logik vorzuweisen. Dies wird sich schwierig darstellen. Priest stellt nicht nur sehr liberal und unter Zuhilfenahme sehr umstrittener philosophischer Ansichten in verschiedensten philosophischen und wissen-

schaftlichen Theorien Widersprüche fest²⁸, sondern verbleibt oft im Bereich der Semantik und des Versprechens einer dazu passenden Beweistheorie²⁹. Auch hier geht es also nicht um eine Interpretation der Position von Priest und Routley sondern um das Verfolgen eines philosophischen/logischen Anliegens.

Diagnose der Probleme der Standard-Logik und Auswege

Die Schwierigkeiten der Standard-Logik, mit inkonsistenten Theorien umzugehen, hängen vor allem mit den sogenannten "Paradoxien der Implikation" zusammen. Die Paradoxien der Implikation sind Theoreme der Standard-Logik, die für ihren konditionalen Junktor gelten, die aber nicht intuitiv unseren Erwartungen an wahre - oder sogar logisch-wahre - konditionale Aussagen entsprechen. Zu diesen Paradoxien der materialen Implikation zählen:³⁰

²⁸ Dazu gehört auch die - nicht-ironisch vorgetragene - These, dass alle großen Systeme der Philosophiegeschichte inkonsistent seien.

²⁹ Unter "Beweistheorie" verstehe ich hier die Angabe eines Kalküls (d.h. die Angabe eines syntaktischen Verfahrens der Manipulation von Zeichenreihen durch Umformungsregeln zur Erzeugung von Theoremen). In *In Contradiction* findet sich beispielsweise *überhaupt* kein Kalkül der parakonsistenten Logik, obwohl Priest nicht nur einen Abschnitt der "Logik" der Parakonsistenz widmet, sondern sich außerdem ständig darauf bezieht, was "die Logik" der Parakonsistenz zeige. Er argumentiert hier ausschließlich vor dem Hintergrund einer semantischen Charakterisierung des parakonsistenten Enthaltenseins, auf die wir in Kapitel 4 zurückkommen. Es kommt aber mindestens genauso darauf an zu wissen, wie entsprechende Argumentationen und Schlüsse einer parakonsistenten Sprache formal aussehen würden. Der Kalkül aus Priest "To Be and Not to Be", auf den er als Beweistheorie auch in (IC) verweist, besitzt gerade keinen konditionalen Junktor! Gerade auf diesen kommt es aber an, wie wir sehen werden. Ebenfalls beschäftigen werden wir uns noch mit den Schwierigkeiten des Systems **LP**, auf das sich Priest in (LT) bezieht.

³⁰ Betrachtet wird hier nur der Fall der materialen Implikation. Die Paradoxien treten aber auch für die modallogische (strikte) Implikation auf, die als $(A \supset B)$ definiert wird. Da in allen *normalen* Modallogiken die Regel der Necessitation ($\vdash A \rightarrow \vdash A$) und das Distributionsaxiom ($\vdash (A \supset B) \supset (\vdash A \supset \vdash B)$) gelten (vgl. Goldblatt, *Logics of Time and Computation*, S.20), lassen sich aus den Paradoxien der materialen Implikation, die als Theoreme der Aussagenlogik Bestandteil dieser Modallogiken sind, Paradoxien

$$(1) \neg A \supset (A \supset B)$$

$$(2) (\neg A \wedge A) \supset B$$

$$(3) A \supset (B \supset A)$$

$$(4) A \supset (B \vee \neg B)$$

wobei (1) und (2) Formen des *ex contradictione quod libet* sind und (3) sowie (4) das *verum ex quod libet sequitur* ausdrücken können. (3) wird sogar oft als Axiom der Standard-Aussagenlogik verwendet.

Das *ex contradictione quod libet* ist Irrelevant³¹, da wir das Vorliegen *irgendeines* Widerspruches nicht als hinreichenden Grund für eine inhaltlich nicht damit zusammenhängende Aussage ansehen. Das *verum ex quod libet sequitur* ist Irrelevant, da es nicht unserem intuitiven Begriff der Implikation entspricht, dass eine wahre Aussage von einer beliebigen anderen Aussage impliziert wird. Mit dem intuitiven Begriff der Implikation erwarten wir einen inhaltlichen Zusammenhang zwischen dem Antezedens und der Konsequenz.

Eine parakonsistente Logik muss nicht alle diese Paradoxien vermeiden, jedoch mindestens alle Versionen des *ex contradictione quod libet*.

Parakonsistente Logiken können ansetzen an den Voraussetzungen einer Ableitung des *ex contradictione quod libet*:

- | | | |
|----|--------------------------------|--------------------------------|
| 1. | A | AE |
| 2. | $\neg A$ | AE |
| 3. | $A \vee B$ | $\vee E, 1$ |
| 4. | B | $\vee B, 2, 3$ |
| 5. | $\neg A \supset B$ | $\supset E, 2, 4$ |
| 6. | $A \supset (\neg A \supset B)$ | $\supset E, 1, 5 \blacksquare$ |

der strikten Implikation ableiten (per Necessitation, dem Distributionsaxiom und der Definition der strikten Implikation).

³¹ "Irrelevant" in dem nun zu erklärenden terminologischen Sinn, der auch in "Relevanter Logik" vorkommt, schreibe ich groß, um es von "irrelevant" zu unterscheiden, wenn ich z.B. behaupte, dass einige Relevante Logiken (philosophisch) irrelevant sind.

Eine Ablehnung des disjunktiven Syllogismus (der $\vee$ -Beseitigung) würde $A, \neg A \vdash B$ blockieren³². Die Schritte (5) und (6) machen nur Schlussprämisse zu Satzprämisse: auch bei Ablehnung der Konditionalisierung ($\supset$ -Einführung) bliebe $A, \neg A \vdash B$ gültig. Es liegt also nahe, den Disjunktiven Syllogismus aufzugeben.

Parakonsistente Logiken können die Versionen des *ex contradictione quod libet* vermeiden, ohne das *verum ex quod libet sequitur* zu vermeiden, indem sie die Negation anders auffassen als in der Standard-Logik - wie wir uns in Kapitel 3.2f. ansehen werden - und so $(\vee B)$ ablehnen.

Parakonsistente Logiken, die die Negation nicht radikal neu auffassen, müssen den konditionalen Junktor neu auffassen, um $(\vee B)$ abzulehnen, da in der Standard-Logik gilt

$$(D) A \supset B \equiv \neg A \vee B$$

so dass bei Beibehaltung dieser Äquivalenz die Ablehnung des Disjunktiven Syllogismus einer Ablehnung des Modus Ponens gleich käme³³, was den elementarsten logischen Intuitionen zuwiderläuft.

Nicht dialethische Ansätze, in denen Widersprüche nie wahr sind und die weder die Negation noch den konditionalen Junktor neufassen, müssen verhindern, dass im Falle, dass in einer Prämissenmenge A und $\neg A$ vorkommen, ein expliziter Widerspruch zustandekommt. Denn sie müssen folgendes Argument (vgl. OP:156f.) verhindern, in dem " $\rightarrow$ " für die strikte Implikation und " $\leftarrow \rightarrow$ " für die strikte Äquivalenz stehen sollen³⁴.

da auch modallogisch gilt:

$$1. A \wedge B \rightarrow B$$

³² Zur Erinnerung: " $A, \neg A \vdash B$ " besagt, dass B aus der Menge der Prämisse $\{A, \neg A\}$ ableitbar ist.

³³ Der Disjunktive Syllogismus besagt ja $A, \neg A \vee B \vdash B$ und setzen wir gemäß der Äquivalenz (D) ein ergibt sich $A, A \supset B \vdash B$, d.h. der Modus Ponens!

³⁴ Diese beiden Symbole werden sonst nirgendwo mehr auftauchen.

gilt:

$$2. \neg B \wedge B \rightarrow B$$

Da nun alle Widersprüche nie wahr sind, gilt:

$$3. \neg B \wedge B \leftarrow \rightarrow \neg A \wedge A$$

also aufgrund der Substitution von Äquivalentem:

$$4. \neg A \wedge A \rightarrow B$$

und das heißt: bei Vorliegen von $(\wedge E)$ würde gelten:

$$4. \neg A, A \vdash B$$

d.h. *ex contradictione quod libet*. Es wäre hier also eine Regel der Konjunktion, $(\wedge E)$, aufzugeben, um Relevant zu bleiben.³⁵

Ansätze der ersten Art (Reinterpretation der Negation) kann man, da sie das *verum ex quod libet sequitur* beibehalten, als "positiv-plus-Ansätze" bezeichnen; solche der zweiten Art (Reinterpretation des konditionalen Junktors) als "Relevante Ansätze"; solche der dritten Art (Reinterpretation der Konjunktion) als "nicht-adjunktive Ansätze"³⁶ (vgl. OP:156f.).

Die geschilderten Varianten der parakonsistenten Logik decken die beweistheoretischen Möglichkeiten, *ex contradictione quod libet* zu

³⁵ Ganz allgemein ergibt sich die Irrelevanz der strikten Implikation in den intuitiv akzeptablen, aber deduktiv sehr starken Modallogiken wie S5 wie folgt (vgl. OP:175): Wenn $\varphi(A, B)$ irgendeine modale Aussage ist mit zwei Teilaussagen A und B sowie einem konditionalen Junktors φ , dann gilt entweder nicht- $(\models \varphi(A, A))$, d.h. Identität (wie wir sie mit " $p \supset p$ " z.B. für die materiale Implikation kennen) bezüglich φ geht fehl oder $\models \varphi(A, A)$. Ist A logisch wahr, dann muss gelten $\models \varphi(A, A)$. B komme in A nicht vor, so gilt trotzdem (z.B. in S5): $\models (A \leftarrow \rightarrow B \vee \neg B)$. Also aufgrund der Substitution von logischen Äquivalenten: $\models \varphi(A, B \vee \neg B)$ - und das ist irrelevant, nämlich von der Art *verum ex quod libet sequitur*, da es keinen Zusammenhang zwischen A und B gibt.

³⁶ Denn $(\wedge E)$ wird oft als "Adjunktionsregel" bezeichnet. Leider wird im Deutschen auch das nicht-ausschließende "oder" manchmal "Adjunktion" genannt (z.B. Kutschera/Breitkopf, *Einführung in die moderne Logik*, S.25ff.).

verhindern ab, denn sie ergaben sich ja aus dem Versuch der Vermeidung seiner Standard-Herleitung.³⁷

Um sie beurteilen und evtl. einige zu verwerfen, benötigen wir über die Zurückweisung des *ex contradictione quod libet* hinausgehende Adäquatheitskriterien für parakonsistente Logiken.

Verschiedene parakonsistente Ansätze

Ausweichen in die Semantik?

Die Trivialität widersprüchlicher Systeme wurde bisher immer als die Ableitbarkeit beliebiger Aussagen verstanden. So genommen ist Inkonsistenz ein Problem der Beweistheorie oder Syntax. Könnte es also vielleicht einen Ausweg aus der Trivialitätsdrohung geben, indem nicht der Kalkül selbst (d.h. die Ableitbarkeitsbeziehung) geändert wird, sondern stattdessen die Semantik ableitbarer Aussagen?

Diesen Weg versuchen Rescher und Brandom in *The Logic of Inconsistency* zu gehen.

Ausgangspunkt sind mögliche Welten. Die aktuelle Welt wird als konsistent angesehen. Doch muss dies nicht für alle Welten gelten. Einige Weltbeschreibungen mögen Widersprüche mit sich bringen. Konsistent kann *über* diese Welten geredet werden.

Der ontologische Status von A ist unabhängig von demjenigen von $\neg A$. Daher können beide vorliegen: A ist der Fall, und $\neg A$ ist der Fall. Nur

³⁷ Wie zu sehen stellt sich das Problem der Parakonsistenz hauptsächlich auf der Ebene der Aussagenlogik (in allen in diesem Kapitel betrachteten Ableitungen spielen Quantoren keine Rolle). Für viele Zusammenhänge reicht es deshalb im Folgenden aus, aussagenlogische Systeme zu betrachten.

$A \wedge \neg A$ ist niemals gegeben. Das "der Fall sein" eines Sachverhaltes verhält sich also - analog gesprochen - "nicht wahrheitsfunktional" (bezüglich der Konjunktion): Auch wenn jeder der beiden Sachverhalte der Fall ist, so ist doch nicht der komplexe Sachverhalt gegeben!

Der Wert von A ergibt sich aus den Werten $[A]_w$ und $[\neg A]_w$, wobei $[\alpha]_w$ den Wert von α in der Welt w bezeichnet. In einer Standard-Welt ist *per definitionem* $A \wedge \neg A$ nie der Fall. Allerdings gibt es Nicht-Standard-Welten. Nicht-Standard-Welten ergeben sich entweder durch die Konjunktion (bzw. den Schnitt) von Welten, wofür gilt:

$$(1) [A]_{w1 \cap w2} = 1 \text{ genau dann, wenn } [A]_{w1} = 1 \text{ und } [A]_{w2} = 1$$

oder durch die Disjunktion (bzw. die Vereinigung) von Welten, wofür gilt:

$$(2) [A]_{w1 \cup w2} = 1 \text{ genau dann, wenn } [A]_{w1} = 1 \text{ oder } [A]_{w2} = 1.$$

Durch die Weltkonjunktion $\cap$ entstehen schematische Welten, wenn A nicht in allen konjungierten Welten der Fall ist (so dass nicht $[A]_{w1 \cap w2} = 1$) und zugleich $\neg A$ nicht in allen konjungierten Welten der Fall ist (so dass nicht $[\neg A]_{w1 \cap w2} = 1$). In einer so entstandenen Welt $w1 \cap w2$ gelten weder A noch $\neg A$, d.h. diese Welt ist bezüglich A unterbestimmt.

Bei der Weltdisjunktion $\cup$ kann es zu widersprüchlichen Bewertungen kommen. Wenn A in einer der disjungierten Welten der Fall ist (z.B. $[A]_{w1} = 1$, so dass $[A]_{w1 \cup w2} = 1$) und zugleich $\neg A$ in einer der disjungierten Welten der Fall ist (z.B. $[\neg A]_{w2} = 1$, so dass $[\neg A]_{w1 \cup w2} = 1$). In der Welt $w1 \cup w2$ liegen dann ein Sachverhalt *und seine Negation* vor.

Rescher und Brandom behalten nun bezüglich von Ableitungen aus den Aussagen, die in einer solchen Welt wahr sind, die Standard-Logik bei! Sie ändern die semantische Relevanz der Ableitungsbeziehung. Insbesondere gilt nicht mehr:

$$(3) A_1 \dots A_n \vdash B \rightarrow \text{wenn } [A_1]_{w1} = 1 \dots [A_n]_{w1} = 1, \text{ dann } [B]_{w1} = 1$$

d.h. eine syntaktische Konklusion aus einer Prämissenmenge muss nicht wahr sein (in einer Welt), obwohl alle Prämissen (in derselben Welt) wahr sind. Stattdessen soll gelten:

$$(4) A_1 \dots A_n \vdash B \rightarrow \text{wenn } [A_1 \wedge \dots \wedge A_n]_{w1} = 1, \text{ dann } [B]_{w1} = 1$$

d.h. wenn die Konjunktion der Prämissen in *einer* Welt wahr ist, so ist auch jede Konsequenz, die aus diesen Prämissen logisch gezogen werden kann, wahr. Im Falle kontradiktorischer Prämissen wäre das Antezedens auf der rechten Seite von " $\rightarrow$ " in (4) falsch (denn eine explizite Kontradiktion soll *nirgendwo* - genauer: zumindest nicht in den Standard-Welten, in denen wir uns überhaupt nur befinden können - wahr sein), so dass (4) nicht zur Anwendung kommt, wir also die Konsequenz einer inkonsistenten Prämissenmenge nicht für wahr halten müssen. Dadurch wird die Trivialität, die ausgehend von einer inkonsistenten Welt droht, vermieden. Dazu muss vorausgesetzt werden, dass wir niemals von einem impliziten Widerspruch (dem Vorliegen von A und $\neg A$) zu einem expliziten Widerspruch ($A \wedge \neg A$) gelangen können. Um dies zu erreichen muss aber die Semantik restringiert werden. So gelten ebenfalls nicht:

$$(5) [A]_{w1} = 1, [B]_{w1} = 1 \rightarrow [A \wedge B]_{w1} = 1$$

$$(6) [\neg A]_{w1} = 1 \rightarrow \text{nicht: } [A]_{w1} = 1$$

Mit (5) wird der Standardwahrheitsfall der Konjunktion aufgegeben. Nur so kann der Übergang vom Vorliegen von A und $\neg A$ zu einer expliziten Kontradiktion (und damit das Erfülltsein des Antezedens der rechten Seite von (4)) blockiert werden. Daraus ergibt sich, dass die semantischen Regeln, die den logischen Regeln normalerweise entsprechen (wie (5) der $(\wedge E)$ entspricht), nicht mehr gültig sind. Ein weiteres Beispiel neben der gerade betrachteten Konjunktionseinführung ist sogar der Modus Ponens: in $w1$ seien A und $\neg B$ der Fall, in $w2$ seien $\neg A$ und $\neg B$ der Fall und $w3 = w1 \cup w2$; dann gelten in $w3$: A und $A \supset B$ (da in

w2 aufgrund der Wahrheit von $\neg A$ auch $\neg A \vee B$ wahr ist), aber B ist nicht wahr in w3, da B in beiden Welten w1 und w2 falsch ist, d.h. die semantische Regel des Modus Ponens hat ein Gegenmodell, d.h. sie ist nicht gültig. Damit gilt *semantisch* auch nicht $A \supset (\neg A \supset B)$, das *ex contradictione quod libet*, ein wünschenswertes Resultat.

Widersprüche lassen sich also ausgehend von inkonsistenten Prämissenmengen (inkonsistenten Welten) nicht semantisch vererben. Wir müssen, so Rescher und Brandom, also beim Auftreten von Widersprüchen, nicht alles für wahr halten:

Gegeben diese Einschränkungen ist es ausdrücklich nicht der Fall, das ein Widerspruch, der irgendwo auftritt, sich überall hin fortpflanzt - selbst dann nicht, wenn wir die klassische Logik beibehalten.³⁸

Diese Konzeption täuscht aber nur eine Lösung des Problems vor. Das beweistheoretische Problem, von dem die parakonsistente Fragestellung ausgeht, wurde nicht gelöst, sondern einfach in die Semantik verschoben. Man bedenke: Da Rescher und Brandom die Standard-Logik nicht aufgegeben haben, gilt auch *ex contradictione quod libet*, d.h. in einer widersprüchlichen Welt *sind alle Aussagen ableitbar*. Es gibt keine Aussage, die sich dort nicht ableiten lässt. Diese Welten sind *absolut inkonsistent*³⁹. D.h. wir haben dort alle Aussagen als Theoreme und *müssen uns fragen*, welche von ihnen wahr sind und welche nicht. Entscheiden im metalogischen Sinne können wir das zwar grundsätzlich (auch in einer Standard-Logik) nicht, hier fehlen jedoch auch die Anhaltspunkte eine Aussage, die wir nicht schon als logisch wahr kennen, für wahr zu halten. Auch muss der vorliegende Widerspruch nicht offensichtlich sein. Wir können zwar die Standard-Logik wie üblich verwenden und interpretieren, wenn unsere Prämissenmenge konsistent ist, aber es gibt interessante Theorien (zumindest alle nicht völlig formalisierten Theorien der Wissenschaften),

³⁸ Rescher/Brandom, *The Logic of Inconsistency*, S.22.

³⁹ Zur Erinnerung: Eine Aussagenmenge ist *einfach* inkonsistent, wenn sie zugleich A und $\neg A$ enthält.

wo wir diesbezüglich nicht sicher sein können.⁴⁰ Und dann können wir nach Rescher und Brandons Ansatz aus dem Umstand, dass wir etwas relativ zu dieser Theorie beweisen können, nicht mehr dazu übergehen, es für wahr zu halten. Wenn diese Funktion der Logik jedoch verloren geht, wieso soll man dann überhaupt noch etwas zu beweisen versuchen? Damit führt sich der Ansatz m.E. selbst ad absurdum.

Darüber hinaus ist die Semantik in einem Masse deviant, das - wie wir bei den beweistheoretischen Ansätzen und dem korrespondierenden Problem der Zurückweisung der entsprechenden syntaktischen Regeln sehen werden – über das zur Sicherung der Nicht-Trivialität Nötige hinausgeht. Bezüglich ihrer stellt sich außerdem das Problem des verstärkten Lügners wieder ein.⁴¹ Der vermeintliche Ausweg in oder über die Semantik hat sich somit als Irrweg erwiesen. Die Lösung muss in der Beweistheorie gesucht werden.

Adäquatheitskriterien zur Beurteilung beweistheoretischer Ansätze

Das Zulassen von Widersprüchen und die Ablehnung des *ex contradictione quod libet* allein reichen nicht aus, um eine Logik als akzeptable parakonsistente Logik unseres Umgehens mit inkonsistenten Theorien (und insbesondere der naiven Semantik) anzusehen. Eine parakonsistente Logik soll die Probleme der Standard-Logik vermeiden, aber dabei möglichst keine neuen produzieren. Insbesondere sollte sie bei der Umgehung der Antinomien nicht neue Implausibilitäten generieren. Eine

⁴⁰ Genaugenommen können wir bei fast allen interessanten Theorien nicht sicher sein. Selbst die Formalisierung - wie im Falle der Mengenlehre vor der Entdeckung von Russells Antinomie - schützt vor verdeckten Widersprüchen nicht.

⁴¹ vgl. Rescher/Brandom, *The Logic of Inconsistency*, S.34.

adäquate parakonsistente Beweistheorie mit einer zugehörigen Semantik sollte mindestens die folgenden Bedingungen erfüllen:

- Das *ex contradictione quod libet* darf nicht gelten. Genauer: es muss in einem inkonsistenten System mindestens eine Aussage geben, die nicht ableitbar ist. Dies nenne ich die "schwache Anti-Trivialitätsbedingung" oder "Bedingung der absoluten Konsistenz".
- Jede Logik muss einen konditionalen Junktor besitzen, für den die Regel des Modus Ponens gilt. Wenn wir irgendetwas mit Logik assoziieren, dann sicher den Modus Ponens. Dies nenne ich die "Modus Ponens-Bedingung".
- Die anderen Junktoren (wie die Negation oder die Disjunktion oder die Konjunktion) dürfen nicht auf eine Weise neu definiert werden, die mindestens so unplausibel ist, wie es die Paradoxien der materialen Implikation für viele Sprachphilosophen sind. Dies nenne ich die "extensionale Bedingung", da sie die wahrheitsfunktionalen Junktoren betrifft.
- Schlußregeln, die wir gewöhnlich mit einem konditionalen Junktor verbinden (wie dessen Transitivität) dürfen nicht ohne Not aufgegeben werden, d.h. solange weniger kontra-intuitive Optionen offenstehen. Dies nenne ich die "Bedingung der minimalen Beschädigung". Diese Bedingung läßt Grade ihrer Erfüllung zu.
- Die naive Semantik und die naive Mengenlehre sollen nicht trivial sein. Dies nenne ich die "starke Anti-Trivialitätsbedingung".

Die starke Anti-Trivialitätsbedingung ergibt sich aus der philosophischen Motivation der parakonsistenten Logik. Sie geht weit über die Bedingung der absoluten Konsistenz hinaus. Dies zeigt die Currys Paradox, das in verschiedenen Formen auftreten kann. Die Formen von Currys Paradox zeigen, dass unter bestimmten Bedingungen (genauer: beim Vorliegen bestimmter logischer Wahrheiten bezüglich des konditionalen Junktors in Verbindung mit einigen üblichen Schlußregeln) die naive Semantik und

die naive Mengenlehre trivial sind. Betrachten wir zwei Versionen von Curry's Paradox:

(I)

(vgl. OP:172f.)⁴²

α sei die Aussage, die behauptet: Wenn diese Aussage wahr ist, dann ist A der Fall. Wobei A eine beliebige Aussage ist. Das heißt:

$$\alpha = "T\alpha \supset A".^{43}$$

In einer semantisch geschlossenen Sprache, die über ihre eigenen Ausdrücke reden kann, muss es eine Aussage wie α geben. Nun gilt aufgrund der Konvention (T):

$$(1) T\alpha \equiv T\alpha \supset A$$

Wenn die Aussage α wahr ist (linke Seite), dann ist α der Fall (rechte Seite).

Eine logische Wahrheit ("Absorption") (z.B. für die materiale oder strikte Implikation) lautet:

$$(2) (A \supset (A \supset B)) \supset (A \supset B)$$

Da nun (1) als Äquivalenz auch eine Implikation ist,

$$(1') T\alpha \supset (T\alpha \supset A)$$

ergibt sich mit (2):

$$(3) T\alpha \supset A$$

daraus ergibt sich per Modus Ponens mit (1)

$$(4) T\alpha$$

⁴² vgl. auch Geach, "On Insolubilia".

⁴³ Als Junktoren werden hier das Konditional und das Bikonditional verwendet. Dies tut der Allgemeinheit des Paradoxes keinen Abbruch, da das Paradox natürlich für alle Junktoren gilt, die in den entsprechenden Positionen vorkommen könnten und für die entsprechende logische Gesetze gelten (also beispielsweise die strikte Implikation und die strikte Äquivalenz). Wahrheit ("T()") wird hier, da von einer semantisch geschlossenen Sprache ausgegangen wird, entweder als Aussagenoperator verstanden (im Sinne von "es ist wahr, dass ()") oder das α im Skopus des Prädikators "T()" (im Sinne von "() ist wahr") ist als Name von α aufzufassen. Das Paradox entsteht in beiden Fällen.

und daraus ergibt sich per Modus Ponens mit (3)

(5) A

wobei A eine beliebige Aussage ist. Es läßt sich also eine beliebige Aussage ableiten: Die naive Semantik ist trivial.

Da wir den Modus Ponens *a/s Regel* für den konditionalen Junktor wohl nicht aufgeben können, ohne unser Verständnis eines konditionalen Junktors überhaupt aufzugeben, müssen wir das Problem in der Verwendung des Absorptionsgesetzes verorten. Die Form (I) von Currys Paradox sagt uns daher:

Eine Logik, die die starke Bedingung der Nicht-Trivialität erfüllen will (d.h. die in der Lage ist, die Logik einer nicht-trivialen universalen Semantik zu sein), und die üblichen Schlußregeln wie den Modus Ponens beibehält, muss das Absorptionsgesetz ungültig machen.

(II)⁴⁴

α sei wieder die Aussage, die behauptet: Wenn diese Aussage wahr ist, dann ist A der Fall. Wobei A eine beliebige Aussage ist. Das heißt:

$$\alpha = "T\alpha \supset A"$$

Dann können wir folgern:

$$1. A \wedge (A \supset B) \supset B$$

ist eine aussagenlogische Tautologie: das Modus Ponens-Theorem.

Durch Einsetzen erhalten wir:

$$2. T\alpha \wedge (T\alpha \supset A) \supset A$$

Durch Einsetzen in die Konvention (T) erhalten wir:

$$3. T\alpha \equiv \alpha$$

also per definitionem von α :

$$4. T\alpha \equiv T\alpha \supset A$$

⁴⁴ vgl. Priest, "Sense, Entailment and *Modus Ponens*", S.431.

Nun erhalten wir per Substitution von Äquivalenten, indem wir mittels (4) in der rechten Seite von (2) einsetzen:

$$5. (T\alpha \supset A) \wedge (T\alpha \supset A) \supset A$$

also per Idempotenz der Konjunktion

$$6. (T\alpha \supset A) \supset A$$

Nun entspricht aber die linke Seite von (6) α und es gilt (3), also per Transitivität:

$$7. T\alpha \supset A$$

Mittels (7) und (4) erhalten wir mit Modus Ponens:

$$8. T\alpha$$

Und aus (8) und (7) mit Modus Ponens:

$$9. A$$

Aus der Kombination einiger aussagenlogischen Tautologien und Schlußprinzipien erhalten wir also jede beliebige Aussage. Welches Glied diese Kette würden wir zuerst aufgeben wollen? Die Transitivität des konditionalen Junktors, der Modus Ponens als Regel und die Substitution von Äquivalenten sind fundamentale Schlußregeln. Wie wir die Idempotenz der Konjunktion ($A \equiv A \wedge A$) aufgeben könnten ist ebenfalls nicht einzusehen. Priest sieht den Schuldigen daher im Modus Ponens-Theorem.⁴⁵

⁴⁵ Priest bezweifelt dabei, dass wir $A \wedge (A \supset B) \supset B$ das Modus Ponens-Theorem nennen sollten. Er plädiert dafür, dass unserem Vorgehen beim Modus Ponens eher entspreche, dass wir ausgehend von einem Konditional $A \supset B$ die Aussage B dann akzeptieren, wenn wir A akzeptieren, also dass dem Modus Ponens das Theorem $(A \supset B) \supset (A \supset B)$ entspreche, eine Instanz des Satzes der Identität (vgl. ebd., S.432f.). Das halte ich nicht für überzeugend: Die Regel des Modus Ponens erwähnt im Antezedens entweder eine Konjunktion von Aussagen oder eine Menge von zwei Aussagen, die wir zugleich als Prämissen haben. Das "A" muss sozusagen vor dem konditionalen Junktor zu stehen kommen. Deshalb trägt das Modus Ponens-Theorem, aufgrund seiner strukturellen Korrespondenz zur Regel des Modus Ponens seinen Namen zu Recht.

Die Bedingungen, die zu Formen von Currys Paradox führen, lassen sich wie folgt systematisieren⁴⁶:

Die Curry Paradoxien lassen sich ableiten, wenn folgende Bedingungen, die jeweils verschieden realisiert werden können, zusammen auftreten:

- I. Konvention (T) (bzw. unbeschränktes Aussonderungssaxiom)
- II. (a) Absorption als Theorem: $(A \supset (A \supset B)) \supset (A \supset B)$
 oder (b) Absorption als Schlußregel: $\vdash (A \supset (A \supset B)) \rightarrow \vdash (A \supset B)$
 oder (c) Modus Ponens als Theorem und (Idempotenz der Konjunktion als Theorem oder $\vdash A \wedge A \supset B \rightarrow \vdash A \supset B$)
- III. (a) Abschluss der Menge der Theoreme unter Ersetzung von Äquivalenten: $\vdash A \equiv B \rightarrow \vdash C \equiv D$, wobei D aus C durch die Ersetzung von A durch B entsteht, oder umgekehrt.
 oder (b) Wenn $\vdash A \equiv B$ und $\vdash (B \supset C)$, dann $\vdash (A \supset C)$
 und: Wenn $\vdash A \equiv B$ und $\vdash B$, dann $\vdash A$
 oder (c) $\vdash A \wedge B \rightarrow \vdash A$ und $\vdash B$
 oder (d) $\vdash A \equiv B$ und $\vdash C \wedge A \supset D \rightarrow \vdash C \wedge B \supset D$
- IV. Abschluss der Menge der Theoreme unter Modus Ponens.

Wollen wir eine nicht-triviale naive Semantik oder naive Mengenlehre⁴⁷ haben, dann müssen wir eine dieser vier Bedingungen fallenlassen. (I) können wir nicht fallenlassen, ohne die naive Semantik oder Mengenlehre fallenzulassen. Wie wir (IV) fallenlassen könnten ist zunächst auch nicht zu sehen: Was soll das denn für ein Konditional sein, für das der Modus Ponens nicht gilt - was heißt denn dann "konditional"? Und die unter (III) angeführten Schlußregeln sind sicherlich fundamentaler als die unter (II) angeführten. Zu verwerfen sind also anscheinend Absorption und Modus Ponens-Theorem. Eine Semantik eines parakonsistenten Kalküls, das die

⁴⁶ vgl. Arruda/da Costa, "On the Relevant Systems P and P^* and Some Related Systems", S.34.

⁴⁷ Für die sich zu den beiden hier betrachteten Formen analoge Curry Paradoxien beweisen lassen.

naive Semantik oder Mengenlehre formalisieren will, muss mindestens in der Lage sein, Absorption und Modus Ponens-Theorem zu verwerfen.⁴⁸ Die Beweistheorie muss gewährleisten, dass sie nicht abgeleitet werden können.

Vorteile und Nachteile einiger parakonsistenter Ansätze

Unter dem Titel "Parakonsistente Logik" verbirgt sich eine Vielzahl logischer Systeme. Viele sind allein aus metalogischen oder formalen Interessen (d.h. Untersuchungen, was für Systeme mit welchen Semantiken sich einführen lassen) motiviert, unabhängig von philosophischen Anliegen im engeren Sinne. Die meisten davon halte ich für philosophisch irrelevant. Die Überlegungen zu den Problemen der Standard-Logik erlauben uns jedoch über die verschiedenen Problemlösungsstrategien einen Blick auf die Typik von parakonsistenten Logiken zu werfen.

⁴⁸ Die Qualifikation "mindestens" ergibt sich, da uns nichts die Vollständigkeit der Bedingungen für Curry Paradoxien garantiert! Smiley versucht ein weiteres solches Paradox zu begründen (vgl. Smiley, "Can Contradictions Be True?", S.29). Er geht aus von der Aussage "Diese Aussage ist nicht beweisbar." Allerdings macht er in dem Beweis sowohl von der Regel der Kontraposition Gebrauch, die sich in Kapitel 4.3 als problematisch herausstellen wird, als auch von einem Prinzip " $\text{Beweisbar}(A \vee B) \wedge \text{Beweisbar}(\neg A) \Rightarrow \text{Beweisbar}(B)$ ", das verdächtig nach dem parakonsistent gewöhnlich ungültigen disjunktiven Syllogismus aussieht. Desweiteren folgert er mit " $\text{Beweisbar}(A) \Rightarrow A$ ", was für ein Standard-Beweisbarkeitsprädikat, aufgrund von Löbs Theorem (vgl. Smullyan, *Gödel's Incompleteness Theorems*, S.109ff.) nicht gilt. Smileys Prämisse " $\text{Beweisbar}(\neg(\text{Beweisbar}(A) \wedge \text{Beweisbar}(\neg A)))$ " muss für einen parakonsistenten Logiker ebenfalls verdächtig klingen. Priest Replik ("Can Contradictions Be True?", S.47ff.) auf Smileys Paradox krankt auch daran, dass Priest " $\text{Beweisbar}(A) \Rightarrow \text{Beweisbar}(\text{Beweisbar}(A))$ " - was für Standard-Beweisbarkeitsprädikate gilt (vgl. Smullyan, *Gödel's Incompleteness Theorems*, S.108f.) - einfach verwirft.

Jaskowski-Systeme

Der sogenannte nicht-adjunktive (nicht-konjunktive) Ansatz geht auf Jaskowski zurück, weswegen auch von "Jaskowski-Systemen" gesprochen wird (vgl. OP:157-62)⁴⁹.

Die Grundidee ist das Referieren oder Zusammenfassen einer sprachlichen Auseinandersetzung oder eines Diskursabschnittes. Die Dinge, die ein Gesprächspartner äußert hält er - von den derivativen Fällen des Täuschens abgesehen - für wahr bzw. stellt er als wahr hin. Den Äußerungen eines Gesprächspartners entspricht dann, insofern sie konsistent sind, eine mögliche Welt, in der das der Fall ist, was er behauptet. Ob dies die wirkliche Welt ist, ist eine andere Frage, aber jedenfalls glaubt der Sprecher das. Wie sieht es nun mit dem Gesamtdiskurs aus? Etwas wird im Gesamtdiskurs für wahr gehalten, wenn mindestens einer der Sprecher es für wahr hält.

Angenommen M wäre ein modalsemantisches Modell mit einer Menge von möglichen Welten, die hier den Sprecheransichten entsprechen sollen. Es läßt sich dann definieren:

$$(1) M \models_d A \leftrightarrow \text{es eine Welt } w \text{ in } M \text{ gibt, so dass } w \models_d A$$

A ist diskursiv wahr in einem Modell M genau dann, wenn es mindestens eine Sprecherwelt in M gibt, in der A wahr ist. Entsprechend ergibt sich dann für die Gültigkeit und die Folgerung:

$$(2) \models_d A \leftrightarrow (\forall M)(M \models_d A)$$

A ist diskursiv gültig, wenn A in allen DiskursModellen diskursiv wahr ist.

$$(3) \Sigma \models_d A \leftrightarrow (\forall M)((\exists B \in \Sigma) \neg M \models_d B \vee M \models_d A)$$

A folgt diskursiv aus einer Prämissenmenge Σ genau dann, wenn entweder nicht alle Prämissen diskursiv wahr sind oder wenn - insbesondere bei Wahrheit der Prämissen - A diskursiv wahr ist. Das setzt allerdings

⁴⁹ vgl. auch Urchs, "Discursive Logic".

nicht voraus, dass alle Prämissen zugleich in einer einzelnen Sprecher-Welt wahr sind!

Gültig sind nach dieser Bestimmung der diskursiven Gültigkeit genau die Aussagen, die auch in der Modallogik S5 gültig sind⁵⁰. Schwieriger liegen die Verhältnisse bezüglich der Folgerungsrelation. Denn es soll in einem Diskurs nicht so sein, dass

$$(4) \{A, \neg A\} \models_d B$$

Aus dem Umstand, dass in einem Diskurs sowohl A als auch $\neg A$ diskursiv wahr sind, und das heißt ja nicht mehr als dass jedes von beiden von mindestens einem Sprecher für wahr gehalten wird, folgt nicht, dass nun die Diskursteilnehmer *alles* für wahr halten. Vielmehr wird es im Diskurs gerade darum gehen nun entweder A oder $\neg A$ zu verwerfen. Die Falschheit von (4) ist also nicht nur eine Forderung an Diskurse, sondern de facto laufen Diskurse nicht so ab. Es gibt also eine Fülle von Gegenbeispielen zu (4). In der diskursiven Logik, den Jaskowski-Systemen, darf (4) also nicht gültig sein. Aufgrund der Weise in der die diskursive Wahrheit mit (1) eingeführt wurde, gibt es nur einen Ausweg: Der Umstand, dass zwei Aussagen diskursiv wahr sind, weil es mindestens zwei Sprecher gibt, die sie für wahr halten, führt nicht dazu, dass sie im Gesamtdiskurs, also von jedem beteiligten Sprecher für wahr gehalten werden. Formal ist dies nichts anderes als die Zurückweisung der Konjunktionseinführung:

$$(5) \text{ Es gilt nicht: } \{A, B\} \models_d (A \wedge B)$$

Damit erhält jedoch den Junktor " $\wedge$ " ein ausgesprochen deviantes Profil. Von der Konjunktion als einem wahrheitsfunktionalen Junktor, der an das umgangssprachliche "und" anknüpft, verlangen wir, dass bei Wahrheit zweier Aussagen auch ihre Konjunktion wahr ist. Somit ermangelt es eigentlich jedes Grundes " $\wedge$ ", das in Jaskowski-Systemen vorkommt, noch

⁵⁰ Das heisst es gilt: $\vdash_d \alpha \leftrightarrow \vdash_{S5} \alpha$ (vgl. Jaskowski, "Propositional Calculus for Contradictory Deductive Systems"). Da hier ja davon ausgegangen wird, dass sich alle Welten (alle Sprecher) wechselseitig zugänglich sind.

als akzeptable Konjunktion zu betrachten. Für die Jaskowski-Konjunktion gibt es *überhaupt* keine rekursive Wahrheitsbedingung ψ derart, dass:

$$(6) M|_d(A \wedge B) \leftrightarrow \psi(M|_d A, M|_d B)$$

d.h. derart, dass die diskursive Wahrheit der Konjunktion allein abhängt von der Wahrheit der Teilaussagen. Denn in faktischen Diskursen wird es immer wieder vorkommen (vgl. OP:159,181) dass zugleich:

$$(7) M|_d B$$

$$(8) M|_d C$$

$$(9) M|_d A$$

$$(10) M|_d(A \wedge B)$$

und

$$(11) \text{Nicht: } M|_d(A \wedge C)$$

d.h. einige wahre Aussagen werden konjungiert, andere aber nicht. Damit kann die Konjunktion in Diskursen aber keine rekursive Bedingung über Wahrheitswerten sein.

Da Costa hat nun vorgeschlagen, eine neue Konjunktion einzuführen, für die Konjunktionseinführung gilt. Als Definition schlägt er vor:

$$(12) A \wedge_d B := \Diamond A \wedge B$$

Warum aber ein Modaloperator wie der der Möglichkeit in die Bedeutung der Konjunktion eingehen soll ist unverständlich. Die unmittelbare Konsequenz davon ist, dass die definitorischen oder aussagenlogischen Beziehungen, die zwischen der Konjunktion und der Disjunktion und der Negation bestehen (z.B. die DeMorgan-Gesetze) nicht mehr gültig sind. Damit weicht diese Konjunktion ebenfalls massiv von unserem Vorverständnis von Konjunktionen ab.

Die Jaskowski-Systeme verstoßen daher gegen die an parakonsistente Logiken gestellte Adäquatheitsbedingung (III), die *extensionale Bedingung*. Außerdem verstoßen sie gegen die Adäquatheitsbedingung (I), die schwache Anti-Trivialitätsbedingung (Ungültigkeit des *ex contradictione quod libet*), insofern in ihnen gilt:

$$(13) \{A \wedge \neg A\} \models_d B$$

da alle aussagenlogischen Tautologien (solange nicht die neue modale Konjunktion eingeführt wurde) gelten. Gerade wegen (13) mußte die Konjunktionseinführung verworfen werden. Im Falle *expliziter* Widersprüche sind Jaskowski-Systeme trivial! In diesen Systemen darf es keine expliziten Widersprüche geben. Damit handelt es sich bei ihnen nicht um die dialetheistische Position (der starken Parakonsistenz). Damit verstoßen sie natürlich auch gegen die starke Anti-Trivialitätsbedingung. Dass keine expliziten Widersprüche zugelassen werden, ist im Übrigen insofern unverständlich als es in Diskursen gerade um das Feststellen von Widersprüchen zwischen den Meinungen verschiedener Sprecher geht und ein dritter Sprecher vielleicht gerade aus den sich widersprechenden Meinungen eine Folgerung ziehen möchte (vgl. OP:161f.).

Bezüglich ihres konditionalen Junktors sind Jaskowski-Systeme zunächst irrelevant, wenn sie die strikte Implikation verwenden, die irrelevant ist. Deshalb hat Joaskowski eine diskursive Implikation eingeführt:

$$(14) A \supset_d B := \Diamond A \supset B$$

Diese diskursive Implikation erfüllt die Modus-Ponens-Bedingung (Adäquatheitsbedingung (II)), da in allen Modallogiken von mindestens der Stärke des Systems **T** gilt:

$$(15) A \supset \Diamond A$$

Allerdings läßt sich auch beweisen (OP:174):

(16) $\models_d A_d$ wobei A_d eine diskursive Aussage ist die aus einer aussagenlogischen Tautologie dadurch entsteht, dass alle Vorkommnisse der materialen Implikation durch die diskursive Implikation ersetzt wurden.

Dieses Ergebnis besagt, dass alle aussagenlogischen Tautologien - insbesondere auch die Paradoxien der materialen Implikation - eine diskursive Entsprechung besitzen. Damit ergibt sich wieder die Irrelevanz

der diskursiven Implikation und ihr Verstoßen gegen die Anti-Trivialitätsbedingungen.

Jaskowski-Systeme sind ein erster Versuch zu einer Logik des intersubjektiven Meinens, wie er in der Standard-Epistemischen-Logik fehlt. Jaskowski-Systeme sind aber in allen für uns interessanten Hinsichten, die sich an einzelnen Meinungssystemen und den zumindest dort wesentlichen Junktoren orientieren, unzureichend. Deshalb werden sie hier nicht weiter betrachtet.

Da Costa-Systeme

Während man bei den im nächsten Abschnitt zu behandelnden Relevanten parakonsistenten Logiken von der "australischen Parakonsistenz" spricht, nennt man die zweite Hauptströmung der parakonsistenten Logik, die von Newton Da Costas (nicht dialethistischen) Ansatz bestimmt wird, die "brasilianische Parakonsistenz".⁵¹ Da Costa-Systeme (vgl. OP:162-67, 175ff.) verwenden eine normale Konjunktion. Sie blockieren das *ex contradictione quod libet* durch ein besonderes Verständnis der Negation. In ihnen gelten als Wahrheitsbedingungen, neben den üblichen Regeln für Konjunktion und Disjunktion:

$$(1) \ v(\neg A)=1 \text{ wenn } v(A)=0$$

$$\text{und } v(A)=1 \text{ wenn } v(\neg\neg A)=1$$

Der erste Teil von (1) entspricht der Bedingung (i), die in Kapitel 1 an eine akzeptable Fassung der Negation gestellt wurde. Der zweite Teil von (1) entspricht der Regel der Negationsbeseitigung, $(\neg B)$ bzw. "Doppelte

⁵¹ Der Ausdruck "parakonsistent" soll auf den brasilianischen Philosophen Miró Quesada zurückgehen (vgl. Arruda, "Aspects of the Historical Development of Paraconsistent Logic"). Die erste Axiomatisierung eines parakonsistenten Kalküls geht, laut Alves ("The First Axiomatization of a Paraconsistent Logic") auf den russischen Logiker

Negation". Aber der zweite Teil von (1) ist schlecht motiviert (vgl. OP:163): Er scheint - wie oben angedeutet - zu sagen, dass $v(\neg\neg A)=1$ meint, dass es nicht der Fall ist, dass $\neg A$. Dies ist zumindest der Grund, der in der Standard-Logik für die Doppelte Negation gilt. In Da Costas Semantik fällt dieser Grund jedoch fort, da sich hier nicht von $v(\neg\neg A)=1$ auf $v(\neg A)=0$ schließen läßt, sondern mit dem ersten Teil von Regel (1) nur umgekehrt von $v(\neg A)=0$ auf $v(\neg\neg A)=1$.

Die zweite Bedingung, die in Kapitel 1 an eine akzeptable Negation gestellt wurde,

(ii) Wenn A wahr ist, ist die Negation von A falsch.

ergäbe sich, wenn

$$(2) \neg(v(A)=0) \equiv v(A)=1.$$

Dies gilt hier indessen nicht, da sowohl A als auch $\neg A$ den Wahrheitswert "1" erhalten können. Dann gilt:

$$(3) v(A)=v(\neg A)=1$$

mit (2) würde sich - außer bei einem Verstoß gegen die Adäquatheitsbedingung, den Modus Ponens zu erhalten⁵² - ergeben:

$$(4) \neg v(A)=0$$

und mit (1) und Kontraposition

$$(5) \neg v(\neg A)=1$$

im Widerspruch zu (3). (2) und damit (ii) oder die Regel der Kontraposition kann also nicht gelten. Wie wir sehen werden gilt diese Richtung der Kontraposition aber in vielen - zumindest den Relevanten - parakonsistenten Logiken. In Da Costa-System gilt allerdings diese Richtung der Kontraposition nicht. Damit wird (ii) also nicht ausgeschlossen. Aber (ii) wird auch nicht gefordert. Und diese Lücke führt zu einem Zusammenbrechen der Rekursivität bezüglich des Bewertens der Negation:

Orlov zurück. Ein Großteil der Arbeiten zur parakonsistenten Logik bezieht sich auf die verschiedensten Varianten von Da Costa-Systemen.

⁵² Für die es nicht wesentlich ist, welche Art von Konditional in (2) auftritt.

Die da Costa-Negation ist überhaupt nicht rekursiv! Denn die Wahrheit von $\neg A$ lässt sich nicht ausgehend von einem Wissen bezüglich A bestimmen. Wenn $v(A)=1$, dann kann nach Regel (1) $v(\neg A)=1$ oder $v(\neg A)=0$ sein. Rekursive Wahrheitsbedingungen sind jedoch unerlässlich für eine Sprache, die von Sprechern, die über endliche Informationen und Sprachregeln verfügen, verstanden werden soll.

Mit der Regel (1) sind weitere Abweichungen vom intuitiven Verständnis der Negation verbunden. Beispielsweise gilt in Da Costa-Systemen nicht der Satz vom Widerspruch, dadurch ist $A \wedge \neg A$ nicht logisch falsch, so dass der mit der Negation verbundene Gegensatz in Da Costa-Systemen, der die Wahrheit von $A \vee \neg A$ verbürgt und nicht verlangt, dass A und $\neg A$ nicht zusammen wahr sein können, dem subkonträren statt dem kontradiktorischen Gegensatz entspricht (vgl. OP:165)⁵³. Dieses subkonträre Verhalten der Negation zeigt sich auch darin, welche logischen Gesetze der Standard-Logik in Da Costa-Systemen nicht mehr gelten, beispielsweise:

$$(6) \{ \neg A \} \models (\neg(A \wedge B)), \{ \neg(A \vee B) \} \models \neg A$$

sowie DeMorgans-Gesetze. Die Da Costa-Negation ist nicht die gewöhnliche Negation, die kontradiktorische Gegensätze bildet, sondern eine subkonträre Negation. Die normale Negation kann auch nicht einfach durch die Hinzunahme von (ii) wieder eingeführt werden. Der erste Teil von (1) und (ii) ergeben:

$$(7) v(\neg A)=1 \text{ genau dann, wenn } v(A)=0.$$

Fügt man aber (7) zu Da Costas Semantik hinzu, die nicht zulässt, dass ein und dieselbe Aussage zwei Werte zugeordnet bekommt⁵⁴, dann ergibt

⁵³ In anderen parakonsistenten Systemen kann zwar auch $A \wedge \neg A$ wahr sein, wenn beide Teilaussagen wahr sind - das ist ja gerade ein Anliegen der parakonsistenten Logik -, da jedoch außerdem der Satz vom Widerspruch gilt, ist $A \wedge \neg A$ zugleich (als Negation einer logisch wahren Aussage) *logisch falsch*, also selbst eine Dialetheia. In ihnen bilden A und $\neg A$ somit einen kontradiktorischen Gegensatz.

⁵⁴ Da Costa lässt zwar zu, dass eine Aussage und ihre Negation denselben Wert bekommen (das war (3)), davon ist aber zu unterscheiden das ein und dieselbe Aussage sowohl den Wert "wahr" als auch den Wert "falsch" erhält.

sich aus der nach (7) zwingenden Falschheit von $A \wedge \neg A$ und der auch von Da Costa verwendeten üblichen Definition der Folgerung

$$(8) \Sigma \models_c A \leftrightarrow (\forall v)((\exists B \in \Sigma) \neg v \models_c B \vee v \models_c A)$$

dass

$$(9) \{A \wedge \neg A\} \models_c B$$

d.h. das System verhielte sich nicht mehr parakonsistent, verstieße also gegen die Adäquatheitsbedingung I, das *ex contradictione quod libet* auszuschließen. Adäquatheitsbedingung I kann indessen nur erfüllt werden auf Kosten der Adäquatheitsbedingung III, des Beibehaltens des intuitiven Verständnisses der extensionalen Junktoren, hier auf Kosten der Negation.

Bekannt geworden sind die Da Costa-Systeme durch eine weitere Konstruktion: Da Costa stellt fest, dass sich alle Aussagen, für die (3) nicht gilt, gemäß der Standard-Logik verhalten. Er wählt für solche Aussagen (der ersten Standard-Stufe⁵⁵) die Notation A^0 und postuliert:

$$(10) \text{ Wenn } B \text{ eine komplexe Aussage mit den Bestandteilen } A_1 \dots A_n \text{ ist} \\ \text{und } \Sigma \models_c A_1^0 \dots A_n^0 \text{ dann } \Sigma \models_c B, \text{ wenn } B \text{ eine Standard-Folgerung aus } \Sigma \text{ ist.}$$

d.h. dass, sofern keine widersprüchlich bewerten Aussagen in einer komplexen Aussage vorkommen, diese sich genau wie in der Standard-Logik verhält. Aussagen, die sich nur aus nicht-widersprüchlichen Aussagen zusammensetzen, haben selbst auch nur Standard-Bewertungen.

Aber zugleich gilt:

$$(11) \{B \wedge \neg B \wedge B^0\} \models_c A$$

d.h. bezüglich eines in Standard-Aussagen auftretenden Widerspruchs gilt *ex contradictione quod libet*. Damit ergibt sich als Problem: In einer

⁵⁵ Da Costa führt nämlich nicht nur eine Stufe, sondern eine Hierarchie von abzählbar unendlich vielen Stufen ein. Dieses Wiedereinführen einer Hierarchie gibt den - nach Kapitel 1 - entschiedenen Vorzug einer parakonsistenten Logik preis: das Umgehen der Schwierigkeiten, semantische Hierarchien zu beherrschen oder zu beschreiben. Da

semantisch geschlossenen Sprache lässt sich eine Aussage konstruieren, die besagt

(12) Aussage (12) ist falsch und standard-bewertet.

Diese Aussage kann nicht wahr sein, soll ein Widerspruch vermieden werden, also muss sie falsch sein. Denn wäre (12) wahr, wäre (12) falsch (d.h. nach (1) $v(\neg(12)=1)$ und standard-bewertet (d.h. $(12)^0$), was dem Antededenz von (11) gleich käme! Ist (12) aber nur falsch, besitzt (12) eine Standard-Bewertung und es lässt sich mit aussagenlogischen Mitteln beweisen:

(13) $(12) \wedge \neg(12) \wedge (12)^0$

und damit nach (11) jede beliebige Aussage A.

Außerdem verstoßen einige Da Costa-Systeme gegen die starke Anti-Trivialitätsbedingung (vgl. OP:175ff.). Da Costa-Systeme werden deshalb hier nicht weiter betrachtet.

Einige Relevanzlogiken

Die Preisgabe der Eigenschaften der Negation und der Konjunktion wäre ein hoher Preis für ein parakonsistentes Logiksystem. Werden Konjunktion und Negation beibehalten, zumindest in dem Sinne, dass $\neg A$ falsch ist, wenn A wahr ist, und $\neg A$ wahr ist, wenn A falsch ist, dann muss man, um das *ex contradictione quod libet* zu vermeiden, entweder den disjunktiven Syllogismus oder die Transitivität des konditionalen Junktors aufgeben.

Systeme, die so verfahren, erfüllen die Adäquatheitsbedigungen (I) und (II), da sie Relevant sind und die elementaren Junktoren nicht uminterpretieren.

Costa-System vereinigen dann nicht das Beste, sondern das Schlechte aus "beiden Welten"(der Logik).

Die Hauptfrage ist hier, wie der konditionale Junktor interpretiert wird. Es kann sich nicht um die materiale Implikation handeln, sondern es muss eine Art von Enthaltensein vorliegen, welche die enthaltende und die enthaltene Aussage in einen inhaltlichen Zusammenhang bringt.

Nach der obigen Ausgangslage gibt es zwei Arten von Systemen: solche ohne disjunktiven Syllogismus und solche ohne Transitivität. Beide Schlußweisen erscheinen gleichermaßen grundlegend und unverzichtbar. Die Transitivität ist insofern fundamentaler, als wir mit dem Schließen die Vorstellung einer Kette aller Konsequenzen, die sich schrittweise ergeben, verbinden. Darauf beruht auch die Idee eines "logischen Abschlusses" einer Aussagenmenge: der logische Abschluss umfaßt nicht nur die primären Konsequenzen, die sich aus den Ausgangsaussagen mit den Schlußregeln ergeben, sondern auch wieder deren Konsequenzen usw., wobei die zuletzt erreichten Aussagen, da sie von Aussagen abhängen, die selber Konsequenzen waren, schließlich auf die Ausgangsmenge zurückverfolgt werden können. All dies erlaubt die Transitivität. Die Systeme, die hier zunächst betrachtet werden und das System **SKP** behalten daher die Transitivität bei und verzichten auf den disjunktiven Syllogismus. Das System von Tennant behält den disjunktiven Syllogismus bei und verzichtet auf die Transitivität. In beiden Fällen verhält es sich jedoch so, dass die jeweils verworfene Schlußregel nur im Falle, dass Antinomien oder Widersprüche unter den Prämissen auftreten, nicht korrekt ist. Im konsistenten Fall gelten sowohl Transitivität als auch disjunktiver Syllogismus.

Relevante Kalküle, die auf den disjunktiven Syllogismus verzichten, sind solche, wie sie u.a. von Anderson/Belnap/Dunn⁵⁶ und Routley vorgeschlagen werden. Die Problematik betrifft das aussagenlogische Schließen, weswegen, da mit der Quantorenlogik keine spezifischen weiteren Relevanzprobleme auftreten, wir uns hier zunächst auf die

Betrachtung aussagenlogischer Systeme beschränken können. Die Syntax (d.h. die Ableitungstechnik oder Beweistheorie), welche die Relevanzkriterien erfüllt, läßt sich angeben, z.B. das System **DL**⁵⁷:

$$A1 \quad A \Rightarrow A$$

Eine Aussage enthält sich selbst.

$$A2 \quad (A \Rightarrow B) \wedge (B \Rightarrow C) \Rightarrow (A \Rightarrow C)$$

Die hier geltende Transitivität.

$$A3 \quad A \wedge B \Rightarrow A$$

$$A4 \quad A \wedge B \Rightarrow B$$

Die beiden Varianten der Konjunktionsbeseitigung.

$$A5 \quad (A \Rightarrow B) \wedge (A \Rightarrow C) \Rightarrow (A \Rightarrow (B \wedge C))$$

Das Enthaltene kann zusammengekommen werden.

$$A6 \quad A \wedge (B \vee C) \Rightarrow (A \wedge B) \vee (A \wedge C)$$

Die Konjunktion kann über die Disjunktion distribuiert werden.

$$A7 \quad \neg \neg A \Rightarrow A$$

Doppelte Negation: Bezüglich der Negation verhält sich der Kalkül DL also klassisch, nicht intuitionistisch.

$$A8 \quad (A \Rightarrow \neg B) \Rightarrow (B \Rightarrow \neg A)$$

Eine Weise, in der Kontraposition möglich ist⁵⁸.

$$A10 \quad A \Rightarrow A \vee B$$

$$A11 \quad B \Rightarrow A \vee B$$

⁵⁶ vgl. dies. *Entailment*, z.B. den Kalkül **R**.

⁵⁷ vgl. Routley/Meyer, "Dialectical Logic, classical logic and the consistency of the world"; vgl. auch Routley, "Dialectical Logic, Semantics and Metamathematics", wo die Axiomatisierung einen unspezifizierten Widerspruch $p^0 \wedge \neg p^0$ verwendet. In der Semantik werden dort Sonderklauseln für die Gültigkeit der Widerspruchsregel und der Transitivität gegeben; vgl. auch: Routley/Loparic, "Relevant Logics Without Replacement", wobei die dort vorgestellten Systeme z.T. auf die Negation verzichten oder aufgrund der Verwendung der materialen Implikation (ebd. S.265f.) das Modus Ponens-Theorem bewähren, also nicht stark anti-trivial sind.

⁵⁸ Junktoren des Enthaltenseins lassen aufgrund der mit ihnen verbundenen aufwendigen Semantik mit ihren Besonderheiten oft keine unbeschränkte Kontraposition zu (s.u.).

Die beiden Varianten der Disjunktionseinführung.

$$A12 \quad (A \Rightarrow C) \wedge (B \Rightarrow C) \Rightarrow ((A \vee B) \Rightarrow C)$$

Ein Enthaltenes erhält man auch aus der Disjunktion von Aussagen, die es enthalten, da eine hierzu reicht.

$$A13 \quad \neg A \wedge \neg B \Rightarrow \neg(A \vee B)$$

$$A14 \quad \neg(A \wedge B) \Rightarrow \neg A \vee \neg B$$

Die Dualität von Konjunktion und Disjunktion, wie sie den Standard-Wahrheitstafeln entspricht.

$$\text{Regel1} \quad A, A \Rightarrow B \rightarrow B \quad (\text{Modus Ponens})$$

$$\text{Regel2} \quad A \Rightarrow B \rightarrow \neg B \Rightarrow \neg A \quad (\text{Kontraposition})$$

Dass die Kontraposition hier als Regel vorhanden ist, aber nicht in dieser Form als Axiom oder Theorem, weist auf eine fundamentale Besonderheit von Kalkülen dieses Typs hin: Die Relation des Ableitens betrifft nur die Wahrheitsvererbung von den Prämissen zur Konklusion. Semantisch besagt sie *nicht mehr*, als dass bei Wahrheit der Prämissen die Konklusion nicht falsch ist. Die Relation des Enthaltenseins, die *in der Sprache* vorkommt, *ist stärker*: hier besteht außerdem irgendeine Art von Zusammenhang inhaltlicher Art zwischen Antezedens und Konsequenz. Daher gilt *nicht*:

$$(DT) \quad \vdash A \Rightarrow B \quad \text{wenn} \quad A \rightarrow B \quad (\text{bzw. } A \vdash B)$$

Das Deduktionstheorem ist für das Enthaltensein nicht gültig! Mit dieser Unterscheidung verbindet sich indessen auch die Chance, in den Regeln die logische Stärke zu bewahren, welche die nötige Beschneidung vom materialen Implizieren zum Enthaltensein entfernen muss. Enthalten und Enthaltensein sollen (vgl. MJ:898ff.) auf einem inhaltlichen Zusammenhang beruhen, der über das Vorkommen extensionaler logischer

Ausdrücke hinausgeht. Bei Erfüllen dieser Forderung können auch nicht die Paradoxien der materialen oder strikten Implikation auftreten⁵⁹.

Das Problem von Logiken des Typs **DL** liegt in der Angabe einer diesem Beweissystem adäquaten und zugleich philosophisch begründbaren Semantik des Enthaltenseins!

Routley und Meyer haben dazu eine modale Semantik entwickelt, der eine dreistellige (gegenüber der ansonsten zweistelligen) Zugänglichkeitsrelation zwischen "Welten" oder "Situationen" zugrundeliegt, wobei unvollständige und unmögliche Welten zugelassen sind. Semantisch entscheidend sind die Modellierung der Negation und des Enthaltenseins:

Die Umkehrung a^* einer Welt a ist die Welt, so dass, wenn $\neg A$ in a der Fall ist, A nicht in a^* der Fall ist.

Eine dreigliedrige Zugänglichkeitsrelation soll die Irrelevanz vermeiden. Dabei ist ein Modell ein Tupel $M = \langle T, O, W, R, *, D, I \rangle$ bestehend aus der realen Welt T , den regulären Welten in O (in denen alle Theoreme gelten, was bei Existenz unmöglicher Welten keine Selbstverständlichkeit mehr ist), der Menge der Welten, einer Zugänglichkeitsrelation R , einer einstelligen Operation $*$ bezüglich der Welten, einem Gegenstandsbereich (für den quantorenlogischen Fall) und einer Interpretationsfunktion. Es soll gelten:

$$T \in O \subseteq W$$

Die aktuelle Welt ist eine reguläre Welt, wobei alle reguläre Welten mögliche Welten sind.

$$D \neq \emptyset$$

Der Individuenbereich ist nicht leer.

⁵⁹ Enthaltensein schafft auch einen Kontext, der stärker ist als konsistente logische Äquivalenz (vgl. MJ:904), d.h. von $\vdash A$ und $\vdash B$ kann man nicht übergehen zu $\vdash A \leftrightarrow B$, ansonsten ergäbe sich über die Äquivalenz von $A \wedge \neg A$ mit $B \wedge \neg B$ wieder das *ex contradictione quod libet*.

(RT) "p" ist wahr in M $\leftrightarrow$ $I(p,T)=1$

In einem Modell ist das wahr, was in der aktuellen Welt wahr ist.

Die Zugänglichkeit zwischen zwei Welten a,b " $\leq$ " wird definiert:

$$(D\leq) a\leq b := (\exists x)(O(x) \wedge R(x,a,b))$$

d.h. b ist von a aus zugänglich, wenn es eine reguläre Welt gibt, zu der a und b in der originären Zugänglichkeitsrelation R stehen, wobei

$$a\leq a \quad (\text{Reflexivität})$$

$$a\leq d \wedge R(d,b,c) \rightarrow R(a,b,c)$$

$$a=a^{**} \quad (\text{so etwas wie doppelte Negation})$$

$$a\leq b \rightarrow b^* \leq a^* \quad (\text{so etwa wie Kontraposition})$$

Der Ausgangspunkt der Bewertungsfunktion ist sowohl im aussagenlogischen wie im prädikatenlogischen Fall wie üblich: Jedem Term wird ein Objekt aus D zugeordnet, jedem n-stelligen Prädikator an jeder Welt eine Menge von n-Tupeln, wobei $a\leq b \wedge I(p,a)=1 \rightarrow I(p,b)=1$ für eine beliebige Aussage "p" und $a\leq b \rightarrow I(P^n,a) \subseteq I(P^n,b)$ für irgendeinen n-stelligen Prädikator P^n . Alle Aussagen halten in einer Welt einen der beiden Wahrheitswerte. Die Wahrheitsbedingungen für komplexe Aussagen sind die üblichen mit Ausnahme der Negationsregel und der Regel für das Enthaltensein. Die Negationsregel lautet:

$$(R\neg) I(\neg p,a)=1 \text{ genau dann, wenn } I(p,a^*)\neq 1$$

Bentham, dem man sicher keine Scheu vor "formalen Tricks" zuschreiben kann, bemerkt hierzu:

In welcher Weise kann ein Widerspruch in dieser Semantik wahr sein? Die Antwort ist einfach: $M\models A \wedge \neg A[a]$ genau dann, wenn $M\models A[a]$ und nicht $M\models A[a^*]$. Bis eine weitere Erklärung der Natur von a^* vorliegt, kann man noch nicht einmal im Ansatz sagen, dies sei mehr als bloß ein formaler Trick.⁶⁰

Die Regel für das Enthaltensein lautet:

$$(R\Rightarrow) \quad I(p\Rightarrow q,a)=1 \leftrightarrow (\forall b,c\in W)(R(a,b,c) \wedge I(p,b)=1 \supset I(q,b)=1)$$

⁶⁰ "What is Dialectical Logic", S.341.

Ein Enthaltensein soll wahr sein, wenn von einer Welt aus die Zugänglichkeit zweier Welten zugänglich ist, so dass die Wahrheit von "p" die von "q" in der als sehend gesehenen Welt *material* impliziert.

Ohne hier weiter auf die Details dieser Semantik (insbesondere des " $\Rightarrow$ ") einzugehen, sei auf das Fazit von Anderson/Belnap/Dunn verwiesen, denen - aufgrund ihres eigenen Interesses an einer Relevanten Semantik - sicherlich keine philosophische Überempfindlichkeit gegen formale ad hoc Konstruktionen unterstellt werden darf. Sie stellen bezüglich dieser Semantik fest: "es gibt keine natürliche Grundlage für sie", "sie selbst hat keine enge Anbindung an unsere natürlichen Vorstellungen."⁶¹

Betrachten wir gegenüber dem Ansatz von Routley u.a. den Ansatz von **Tennant**. Tennant verwirft nicht den disjunktiven Syllogismus, sondern zwangsläufig die Transitivität. Dabei geht er nicht von einem modalen Begriff des Enthaltenseins aus. Die mit diesem verbundenen Schwierigkeiten seiner semantischen Explikation kann er somit vermeiden. Er entwickelt seinen Kalkül der Relevanzlogik, indem er an einen nicht Relevanten Kalkül des Natürlichen Schließens weitere Anforderungen stellt. Diese Anforderungen modifizieren so unser Umgehen mit der *materialen* Implikation, wie es scheint. Dass Tennant dabei von einem intuitionistischen Kalkül ausgeht ist für die Frage, die uns hier beschäftigt, nicht weiter wichtig.

Der Kalkül besteht aus folgenden Regeln:⁶²

⁶¹ Anderson/Belnap/Dunn, *Entailment*, II, S.163 und S.164.

⁶² Da es sich um einen Kalkül des Natürlichen Schließens handelt, gibt es keine Axiome. Die Bedeutung und logische Kraft der Junktoren findet sich allein in den Regeln wieder, die auf beliebige Prämissenmengen (also nicht bloß logisch wahre Axiome oder Theoreme) angewendet werden können. Vgl. zu den Regeln: Tennant, *Anti-Realism and Logic*, S.258f.. Die Regeln werden in zweidimensionalem Format dargestellt: der Strich "_" verweist auf das Einführen einer Annahme oder das Ziehen einer Konsequenz; Punkte unterhalb einer Aussage verweisen auf eine beliebige weitere Ableitung, die von dieser Aussage ausgeht; Aussagen, die oberhalb eines Striches *nebeneinanderstehen*, kommen in verschiedenen Tei ableitungen vor, die mittels der Regel, auf die sich der betreffende Strich bezieht, in eine Ableitung der Konsequenz unter dem Strich überführt werden.

R1: $\neg$ E (Negationseinführung)

<u> </u> (i)	eine gemachte Annahme A führt in
A	einer Ableitung zu einem Widerspruch
.	
.	
$\perp$	daraus folgt, bei Entlastung der Annahme,
	ihre
<u> </u> (i)	Negation
$\neg A$	

R2: $\neg$ B (Negationsbeseitigung)

<u>A $\neg A$</u>	aus zwei Teildableitungen einer
$\perp$	Aussage und ihrer Negation folgt der Wider-
	spruch

R3: $\wedge$ E (Konjunktionseinführung)

<u>A B</u>
$A \wedge B$

R4: $\wedge$ B (Konjunktionsbeseitigung)

<u>A $\wedge$ B</u>	<u>A $\wedge$ B</u>
A	B

R5: $\vee$ E (Disjunktionseinführung)

<u>A</u>	<u>B</u>
$A \vee B$	$A \vee B$

R6: $\vee$ B (Disjunktionsbeseitigung)("allgemeines Dilemma")

<u> </u> (i)	<u> </u> (i)	aus einer Disjunktion
A	B	kann man auf eine Aussage
.	.	schließen, wenn sie aus
.	.	beiden Disjunkten
<u>A $\vee$ B</u>	<u>$\perp$/C</u>	<u>$\perp$/C</u> (i)
		herleitbar ist
	$\perp$ /C	

(wobei die Notation " $\perp/C$ " besagt: die Konklusion einer der Teildableitungen kann als Gesamtkonklusion übernommen werden, auch wenn die andere Teildableitung zu " $\perp$ " führt.)

R7: $\supset E$ (Implikationseinführung) ("Konditionalisierung")

$_ (i)$

A aus einer Schlußprämisse kann eine
 . "Satzprämisse" (ein Antezedens eines
 . Konditionals) gemacht werden

B $_ (i)$

$A \supset B$

R8: $\supset B$ (Implikationsbeseitigung) ("Modus Ponens")

A $A \supset B$

B

R9: $\forall E$ (Alleinführung) ("Universelle Generalisierung")

.

.

P(a)

$(\forall x)P(x)$

wobei "a" in keiner Annahme, von der "P(a)" abhängt,
 und nicht in der Konklusion vorkommt.⁶³

R10: $\forall B$ (Allbeseitigung) ("Universelle Spezialisierung")

$(\forall x)P(x)$

P(a)

R11: $\exists E$ (Existenzeinführung) ("Existentielle Generalisierung")

P(a)

$(\exists x)P(x)$

⁶³ Hinzufügen muss man noch das Verbot der Mehrfachmarkierung (dass "a" durch existentielle Spezialisierung eingeführt wurde) und der zirkulären Markierungsabhängigkeit (vgl. Essler/Martinez, *Grundzüge der Logik*, S.194).

R12: $\exists B$ (Existenzbeseitigung)("Existentielle Spezialisierung")

$$\begin{array}{c}
 \text{---(i)} \\
 P(a) \\
 \cdot \\
 \cdot \\
 \hline
 (\exists x)P(x) \quad B \text{---(i)} \\
 B
 \end{array}$$

wobei "a" weder in $(\exists x)P(x)$, B oder einer Annahme außer

$P(a)$, von der das obere Vorkommenis von B abhängt, vorkommt.⁶⁴

Um nun mit diesen Regeln einen relevanten Kalkül zu erreichen werden an die Entlastungen von Prämissen, die in den Regeln immer mit den Indices "(i)" angezeigt wurden, weitere Anforderungen gestellt⁶⁵:

1. In allen Regeln, in denen die Entlastung einer Annahme vorkommt, ist die Entlastung obligatorisch.

Damit wird z.B. verhindert:

$$\begin{array}{c}
 \underline{B} \\
 A \supset B
 \end{array}$$

wo die Konklusion ein Irrelevantes Antezedens aufweist (verum ex quod libet sequitur).

2. Eine Ableitung muss in Normalform sein. Sie ist in Normalform, wenn kein Satzvorkommenis in ihr zugleich die Konklusion einer Einführungsregel und die Hauptprämisse einer Beseitigungsregel bezüglich desselben logischen Ausdrucks ist, wobei Konklusionen von Teableitungen, die dort bezüglich eines logischen Ausdrucks einführend sind, auch als Hauptkonklusion diesbezüglich einführend sind.

⁶⁴ Mit den entsprechenden Ergänzungen wie bei $\forall E$.

⁶⁵ vgl. Tennant, *Anti-Realism and Logic*, S.257.

Mit dieser Normalformbedingung wird die Ableitung des *ex contradictione quod libet* verhindert. Diese sähe so aus:

$$\begin{array}{c}
 \text{---(ii)} \qquad \qquad \text{---(ii)} \\
 \hline
 \text{A} \wedge \neg \text{A} \qquad \neg\text{(i)} \quad \text{A} \wedge \neg \text{A} \qquad \neg\text{(i)} \\
 \hline
 \text{A} \qquad \text{A} \quad \neg \text{A} \qquad \text{B} \\
 \hline
 \text{A} \vee \text{B} \qquad \perp \qquad \text{B} \text{---(i)} \\
 \hline
 \text{B}
 \end{array}$$

Aus der Kontradiktion (der Annahme (ii)) wird links A abgeleitet nach ($\wedge$ B) und in der Mitte $\neg$ A, links dient A als Prämisse in ($\vee$ E) und in der Mitte wird auf die Annahme A und auf das gewonnene $\neg$ A ($\neg$ B) angewendet. B bleibt links als Annahme erhalten. Auf $A \vee B$, $\perp$ und B kann dann ($\vee$ B) angewendet werden. Dadurch werden die markierten Annahmen entlastet, und es gilt:

$$(*) \quad \text{A} \wedge \neg \text{A} \vdash \text{B}$$

Das *ex contradictione quod libet* gilt intuitionistisch. Diese Ableitung gilt aber nicht Relevant, da sie der Normalformbedingung nicht genügt: $A \vee B$ kommt sowohl als Konklusion der Disjunktionseinführung vor (links) als auch als Hauptprämisse in der Disjunktionsbeseitigung (im Gesamtschluß). Daher ist diese Ableitung Irrelevant.

In dem System gilt auch der Disjunktive Syllogismus:

$$(T1) \quad \text{A} \vee \text{B}, \neg \text{A} \vdash \text{B}$$

Beweis:

$$\begin{array}{c}
 \neg\text{(i)} \quad \neg\text{(ii)} \qquad \neg\text{(i)} \\
 \hline
 \text{A} \quad \neg \text{A} \qquad \text{B} \\
 \hline
 \text{---(iii)} \\
 \hline
 \text{A} \vee \text{B} \qquad \perp \qquad \text{B} \text{---(i)} \\
 \hline
 \text{B}
 \end{array}$$

Entlastet werden hier nach ($\vee B$) die beiden vorübergehenden Annahmen (i). In der Mitte ergibt sich aus der Annahme (ii) nach Negationsbeseitigung das $\perp$.

Mit dem Disjunktiven Syllogismus oder mit ($\vee B$) kann man auch beweisen:

$$(T2) \quad \neg A \vee B \supset (A \supset B)$$

Beweis:

$$\begin{array}{c}
 \begin{array}{ccc}
 \neg(i) & \neg(ii) & \neg(i) \\
 \hline
 \neg A & A & B
 \end{array} \\
 \hline
 \neg A \vee B & & \perp \\
 \hline
 & & B \quad (i) \\
 \hline
 & B & (ii) \\
 \hline
 & A \supset B &
 \end{array}$$

(Die zweite Annahme, A, wird hier mit ($\supset E$) entlastet.)

Die Umkehrung von (T2) gilt allerdings nicht, d.h. es gilt nicht

$$(**) (A \supset B) \supset \neg A \vee B$$

Denn mit Disjunktionsbeseitigung kann man ausgehend von $A \vee A$ beweisen, dass

$$(T3) \quad A \vee A \supset A$$

und nach Disjunktionseinführung kann man beweisen:

$$(T4) \quad A \supset A \vee A$$

also mit der Definition des Bikonditionals zusammen:

$$(T5) \quad A \equiv A \vee A$$

und dann mit der Substitution von Äquivalenten auf (T5) selbst angewendet, wieder der Definition des Bikonditionals und der darauf angewandten Konjunktionsbeseitigung:

$$(T6) \quad A \supset A$$

Würde nun (**) gelten, könnten wir aus (T6) ableiten

$$(***) \quad A \vee \neg A$$

was in Tennants Kalkül, da er intuitionistisch ist, nicht gilt. Insofern unterscheidet sich das " $\supset$ " in diesem Kalkül vom Standard-Konditional. Auf diese Weise unterscheidet sich das intuitionistische Konditional aber immer. Schwerwiegender ist das Versagen der Transitivität für

" $\supset$ ". Wir haben als Theoreme:

$$(T7) \quad A \supset A \vee B$$

(nach Disjunktionseinführung)

und

$$(T8) \quad A \vee B \supset (\neg A \supset B)$$

(nach Disjunktivem Syllogismus). Per Transitivität müßten wir also

$$(***) \quad A \supset (\neg A \supset B)$$

erhalten, haben aber schon gesehen, dass (*), wovon (***) nur die durch Konditionalisierung gewonnene Form ist, nicht gültig ist (aufgrund der Normalformbedingung). Deshalb kann keine Transitivität vorliegen, wenn der Kalkül korrekt ist.

Aber ist der Kalkül korrekt? Darauf gibt Tennant keine Antwort, weil er keine Semantik für diesen Kalkül angibt! Ohne eine Semantik für das Konditional läßt sich aber weder die Korrektheit noch die logische Vollständigkeit des Kalküls beurteilen. Die Semantik, die Tennant für den intuitionistischen Kalkül des Natürlichen Schließens gibt⁶⁶, müßte so modifiziert werden, dass die Zusatzbedingungen (1.) und (2.) eine semantische Entsprechung erhalten.

Die intuitionistische Wahrheitsbedingung für das Konditional besagt:

$$(I\supset) \quad i \models A \supset B \text{ genau dann, wenn } (\forall j \geq i)(j \models A \rightarrow j \models B)$$

Die Semantik bezieht sich dabei auf Etappen (i, j usw.) eines Begründungsprozesses. Zu einer Etappe (zu einem Zeitpunkt des Begründungsprozesses) läßt sich ein Konditional behaupten - von Wahrheit (der bivalenten Lesart von " $\models$ ") muß ja zu Behauptbarkeit im Intuitionismus über-

⁶⁶ vgl. Tennant, *Natural Logic*, S.106ff., 132ff.

gegangen werden - genau dann, wenn sich für alle an diesen Begründungsstand anschließenden Etappen des Begründungsprozesses B behaupten läßt, im Falle, dass sich A behaupten läßt. Die Behauptbarkeit einer Aussage bezieht sich dabei auf kanonische Beweisverfahren, die sich auf schon Bewiesenes zurückbeziehen. Ein kanonischer Beweis für ein Konditional $A \supset B$ ist eine *Methode*, die darin besteht, dass sie angewendet auf einen Beweis von A einen Beweis von B liefert⁶⁷. Diesen Begriff müßte man nun für einen Relevanten Kalkül so modifizieren, dass - gemäß Bedingung (1.) - in der Methode, die uns einen Beweis von B liefert, auch auf A *zurückgegriffen* wird *und* ansonsten der Beweis von B nicht zustande käme, dass A also ein wesentlicher Grund für die Behauptbarkeit von B ist. Semantisch müßten wir so etwas haben wie folgende Zusammenhänge: $\Sigma \models A \supset B$, das Konditional ist *wahr* relativ zu einer möglicherweise leeren Menge von Aussagen Σ wenn bezüglich der *in der intuitionistischen Semantik* erwähnten Begründungsrelation " $\vdash_s$ " gilt $\Sigma, A \vdash_s B$, so dass die Hinzunahme von A zu Σ die Begründung von B erlaubt, während nicht $\Sigma \vdash_s B$, Σ alleine also nicht in der Lage ist, B zu begründen.

Eine entsprechende Modifikation und ihrer Übersetzung in eine semantische Regel wie (I $\supset$) bedürfte es für die Bedingung (2.). Beides fehlt.

Aber selbst wenn es gelänge, eine solche Semantik zu liefern, so leidet Tennants Kalkül an einem anderen Defekt: Er erfüllt nicht die starke Anti-Trivialitätsforderung. Relativ zu diesem Kalkül sind die naive Mengenlehre und die naive Semantik trivial, da sich in ihm z.B. das Absorptionsgesetz beweisen läßt.

$$(T9) (A \supset (A \supset B)) \supset (A \supset B)$$

Beweis:

⁶⁷ vgl. ebd. S.108

$$\begin{array}{c}
 \frac{}{}(i) \quad \frac{}{}(ii) \\
 \frac{A \supset (A \supset B)}{A \supset B} \quad \frac{A}{A} \quad \frac{}{}(ii) \\
 \frac{A \supset B}{A \supset B} \quad \frac{A}{A} \\
 \frac{B}{B} (ii) \\
 \frac{A \supset B}{A \supset B} (i) \\
 (A \supset (A \supset B)) \supset (A \supset B)
 \end{array}$$

Im Beweis der Absorption wird die Prämisse (ii) zweimal verwendet. Doch Tennants System läßt dies zu. Würde man dies ad hoc verbieten - was semantisch selbstverständlich zu begründen wäre -, bliebe immer noch der Beweis des Modus Ponens-Axioms

(T10) $A \wedge (A \supset B) \supset B$

Beweis:

$$\begin{array}{c}
 \frac{}{}(i) \quad \frac{}{}(i) \\
 \frac{A \wedge (A \supset B)}{A} \quad \frac{A \wedge (A \supset B)}{A \supset B} \\
 \frac{A}{A} \quad \frac{A \supset B}{A \supset B} \\
 \frac{B}{B} (i) \\
 A \wedge (A \supset B) \supset B
 \end{array}$$

Tennants System verstößt also gegen die starke Anti-Trivialitätsbedingung.

Trotzdem sind Systeme, die gegen die starke Anti-Trivialitätsbedingung verstoßen von Interesse. Im Sinne des schwachen parakonsistenten Ansatzes soll die parakonsistente Logik auch in der Lage sein, Theorien, die im Moment inkonsistent sind oder es waren, zu behandeln, auch wenn die Inkonsistenz logisch vermeidbar gewesen wäre. Des Weiteren mag es philosophisch zwingende Widersprüche geben, wobei die Theorien, in denen sie auftreten, nicht unter Currys Paradox fallen.

Jenseits der naiven Semantik und Mengenlehre kann also ein parakonsistenter Kalkül von Interesse sein, der gegen die starke Anti-Trivialitätsbedingung verstößt. Dies könnte – unabhängig davon, ob dies Tennants eigenen philosophischen Ansichten und seiner *intendierten* Semantik entspräche – Tennants Kalkül sein, der die Extensionalitätsbedingung und die Modus Ponens-Bedingung erfüllt, wenn er mit einer Semantik versehen werden kann.

Verschiedene weitere Ansätze

Sieht man von den Adäquatheitsbedingungen ab und stellt keine philosophischen Relevanzbedingungen auf, so lassen sich beliebig Kalküle formulieren.

Allerdings verstoßen sie dann auch gegen eine der Adäquatheitsbedingungen.

Einige Beispiele:

- Popov formuliert einige parakonsistente Sequenzenkalküle⁶⁸. Ihnen fehlt es sowohl an semantischer Motivation als auch an Regeln (oder Analoga dazu) wie Konjunktionsbeseitigung, die wir nach der Extensionalitätsbedingung erwarten.
 - Bei Schotch und Jennings Art von Sequenzenkalkül⁶⁹ fehlt nicht nur die Konjunktionseinführung, sondern außerdem (!) der Modus Ponens.
 - Meyer and Slaney⁷⁰ führen eine neue Negation ein ("relativierte Negation"), die ihrer algebraischen Semantik adäquat sein soll.
- usw.

⁶⁸ vgl. Popov, "Paraconsistent Sequential Calculi"

⁶⁹ vgl. Schotch/Jennings, "On Detonating".

⁷⁰ vgl. Meyer/Slaney, "Abelian Logic (from A to Z)".

Der stark anti-triviale Relevante Ansatz

Jetzt sollen Systeme vorgestellt werden, die sich dadurch auszeichnen, auch die starke Anti-Trivialitätsbedingung zu erfüllen. Bei ihnen allein kann es sich um die gesuchte Logik der natürlichen Semantik handeln.

Allerdings hat - soweit mir bekannt - niemand einen solchen Kalkül 1.ter Stufe mit einer philosophisch nachvollziehbaren und akzeptablen Semantik bis jetzt vorgelegt. Dieser negative Umstand ergibt sich selbstverständlich aus meinen Standards dessen, was als Semantik philosophisch nachvollziehbar ist. Routley sieht dies – natürlich – ganz anders.

Insbesondere Priest, der neben seinen Büchern viele Aufsätze publiziert hat, hat, obwohl er ständig von der parakonsistenten "Logik" spricht, keinen solchen Kalkül vorgelegt, an dem er dann festgehalten hätte. Vielmehr legte er mehrere Vorschläge vor, die Mängel aufweisen. Außerdem verweist er in Repliken auf Kritiker der parakonsistenten Logik oft darauf, das bestimmte von den Kritikern verwendete Schlußregeln in seiner parakonsistenten Logik nicht gelten würden. So lehnt er gelegentlich das aussagenlogische Gesetz der Importation oder die Transitivität des konditionalen Junktors ab⁷¹, obwohl sie im System von "Sense, Entailment and Modus Ponens" gilt⁷².

Im Folgenden werde ich insbesondere zwei Kalküle diskutieren, eines greife ich in seinem aussagenlogischen Teil mehr oder weniger direkt von Priest auf, vor allem dem anderen ist jedoch erst (u.a. aus verschiedenen Bruchstücken von Ansätzen bei Priest) eine Semantik zu verschaffen. Den Quantifikationsteil übernehme ich zunächst aus der Standard-Logik,

⁷¹ Letzteres z.B. in "Can Contradictions Be True?", S.47.

⁷² Die Kritiken der Parakonsistenten Logik (wie die von Everett und Smiley, s.u.), die Priest u.a. mit der Zurückweisung bestimmter Schlußregeln erwidert, kann man m.E. – ohne, dass ich dies jeweils im Detail zeige - mit den anderen Erwideroptionsen, die Priest vorbringt, zurückweisen. Mit dem Zulassen von Widersprüchen geht eine Beschränkung der Beweistheorie einher. Jede weitere Beschränkung der Beweistheorie sollte gemäß der Bedingung der minimalen Beschädigung vermieden werden.

da die Weise der Quantifikation nicht relevant ist für die Fragen der Parakonsistenz im engeren Sinne. Hätten wir einmal eine parakonsistente Aussagenlogik ließe sich die Frage nach der angemessenen quantifikationalen Logik (ob Freie Logik oder nicht usw.) auch auf die Parakonsistente Logik beziehen. Im Moment geht es aber um die für die Parakonsistente Logik grundsätzlichere Auseinandersetzung um ihre Durchführbarkeit überhaupt.

Im Gegensatz zu (IC) und "The Logic of Paradoxes", deren jeweilige Ansätze von Semantik ich referiere, muss - wie Priest selbst festgestellt hat - die Interpretation nicht von einer Funktion, sondern von einer Interpretationsrelation vorgenommen werden, um den *Hyper-Widersprüchen* zu entgehen.

Starke Anti-Trivialität ohne Modus Ponens!

Im folgenden soll ein erster Kalkül vorgestellt werden, der die starke Anti-Trivialitätsbedingung erfüllt und so zur Behandlung der naiven Semantik geeignet wäre. Es handelt sich um eine quantorenlogische Erweiterung des Kalkül **LP**, den Priest in dem Aufsatz "The Logic of Paradox" semantisch charakterisiert.

Die aussagenlogischen Junkoren werden dort extensional aufgefaßt. Eine Bewertung v ordnet einer Aussage A eine nichtleere Teilmenge zu aus der Menge mit den Wahrheitswerten 1 und 0 als Elementen. Das heißt entweder $\{1\}$ ("wahr") oder $\{0\}$ ("falsch") oder $\{0,1\}$ ("antinomisch")⁷³. Es gelten als semantische Regeln u.a.:

$$(LP1) \text{ (i) } 1 \in v(\neg A) \leftrightarrow 0 \in v(A) \text{ (ii) } 0 \in v(\neg A) \leftrightarrow 1 \in v(A)$$

$$(LP2) \text{ (i) } 1 \in v(A \wedge B) \leftrightarrow 1 \in v(A) \text{ und } 1 \in v(B)$$

⁷³ Genaugenommen muss man - wie schon erwähnt - diese Bewertung, um das Problem der Hyperwidersprüche zu vermeiden, als Relation (nicht wie hier als

$$(ii) 0 \in v(A \wedge B) \leftrightarrow 0 \in v(A) \text{ oder } 0 \in v(B)$$

An diesen Wahrheitsbedingungen sieht man, dass **LP** die Extensionalitätsbedingung erfüllt. Der erste Teil von (LP1) ist die Standard-Wahrheitsbedingung für die Negation. Mit dem zweiten Teil wird sie lediglich bezüglich des Falls antinomischer Aussagen erweitert. Die LP-Negation ist also eine Erweiterung der Standard-Negation. Die wesentlichen Bedingungen, an denen Da Costa-Systeme scheiterten, werden von **LP** erfüllt. Um an unsere beiden Klauseln (i) und (ii) bezüglich der Negation anzuknüpfen, könnten wir auch sagen:

(LP3) (i') $\neg A$ ist mindestens wahr, wenn A mindestens falsch ist.

(ii') $\neg A$ ist mindestens falsch, wenn A mindestens wahr ist.

Genauso entspricht der erste Teil von (LP2) den Standard-Wahrheitsbedingungen der Konjunktion, an deren Nichterfüllung Jaskowski-Systeme scheiterten. Auch diese Standard-Bedingung wird ergänzt um eine Zusatzbedingung, die sich durch das Vorliegen antinomischer Aussagen ergibt. Entsprechende Regeln für die Disjunktion und das Konditional lassen sich mit den beiden erwähnten Regeln am besten in LP-Wahrheitstafeln darstellen:

A	B	$\neg A$	$A \wedge B$	$A \vee B$	$A \supset B$
1	1	0	1	1	1
1	0	0	0	1	0
1	0,1	0	0,1	1	0,1
0	1	1	0	1	1
0	0	1	0	0	1
0	0,1	1	0	0,1	1
0,1	1	0,1	0,1	1	1
0,1	0	0,1	0	0,1	0,1
0,1	0,1	0,1	0,1	0,1	0,1

Funktion) auffassen. Aus Einfachheitsgründen bleibe ich hier bei der funktionalen

$A \equiv B$ kann man wie üblich definieren als $(A \supset B) \wedge (B \supset A)$; die materiale Äquivalenz ist dann bloß wahr (d.h. nicht zugleich falsch), wenn A und B beide bloß wahr oder beide bloß falsch sind - sie stimmt also wieder im Bereich der Standard-Werte mit der Standard-Semantik überein -, sie ist entsprechend dann bloß falsch, wenn eine von den beiden Aussagen bloß wahr und die andere bloß falsch ist, und in allen anderen Fällen (d.h. wenn eine der beiden Aussagen antinomisch ist) ist die materiale Äquivalenz antinomisch.

Zu diesen Tafeln einige Erläuterungen:

Die Negation einer Antinomie ist eine Antinomie.

Eine Konjunktion ist bloß wahr genau dann, wenn die beiden Konjunkte wahr sind. Eine Konjunktion mit mindestens einer Antinomie ist antinomisch.

Die Disjunktion ist ausschließlich falsch, wenn beide Disjunkte falsch sind. Die ausschließliche Wahrheit eines Disjunks genügt, um die Disjunktion wahr zu machen.

Das Konditional ist nur falsch, wenn bei bloßer Wahrheit des Vordersatzes der Hintersatz ausschließlich falsch ist. Auch dies entspricht der Standard-Wahrheitsbedingung. Ist der Vordersatz ausschließlich falsch, ist das Konditional ausschließlich wahr⁷⁴. Ist der Vordersatz antinomisch, ist das Konditional nicht immer wahr, sondern außer im Fall, dass der Hintersatz wahr ist, auch antinomisch. Dass das Konditional bei rein wahrem Hintersatz auch rein wahr ist, ist ungefährlich, da wir so nur bei einem Hintersatz, der eh schon wahr ist, bleiben.

Insbesondere gilt in **LP** die Äquivalenz

$$(LPT1) \quad A \supset B \equiv \neg A \vee B$$

die in Tennants System nicht galt.

Darstellung.

⁷⁴ Daraus ergibt sich, dass einige Paradoxien der materialen Implikation auch in **LP** gelten. Daneben auch z.B. auch das *verum ex quod libet sequitur*.

Die Wahrheitstafeln in **LP** sehen dabei wie dreiwertige Wahrheitstafeln aus. Würden wir "0,1" durch "unbestimmt" ersetzen, erhielten wir die Wahrheitstafeln von Lukasiewiczs dreiwertiger Logik:

A	B	$\neg A$	$A \wedge B$	$A \vee B$	$A \supset B$
1	1	0	1	1	1
1	0	0	0	1	0
1	unbest.	0	unbest.	1	unbest.
0	1	1	0	1	1
0	0	1	0	0	1
0	unbest.	1	0	unbest.	1
unbest.	1	unbest.	unbest.	1	1
unbest.	0	unbest.	0	unbest.	unbest.
unbest.	unbest.	unbest.	unbest.	unbest.	unbest.

Nun ist Lukasiewiczs Logik aber eine Beschränkung der Standard-Logik und **LP** soll Standard-Theoreme wie den Satz vom Ausgeschlossenen Dritten beibehalten. Wie ist das möglich? Nun, eine logisch wahre Aussage soll immer wahr sein. Eine Aussage ist in **LP** allerdings auch dann wahr, wenn sie antinomisch ist, da Antinomien zugleich wahr und falsch sind. Der ausgezeichnete von den Wahrheitswerten ist also nicht wie sowohl in der Standard- als auch in der dreiwertigen Logik üblich das Wahre (verstanden als rein Wahres), sondern das Wahre, sei es auch in einer Antinomie mit dem Falschen verknüpft. Der ausgezeichnete Wahrheitswert ist derjenige, dessen interpretationsübergreifendes Zukommen Gültigkeit ausmacht. Formal ausgedrückt für **LP**:

$$(LPG) \quad \models_{LP} A \leftrightarrow (\forall v)(1 \in v(A))$$

Eine Aussage A ist gültig in **LP** genau dann, wenn sie von allen Bewertungen mindestens als wahr bewertet wird. Entsprechend lässt sich die Folgerung definieren:

$$(LPF) \quad \Sigma \models_{LP} A \leftrightarrow (\forall v)(1 \in v(A) \vee (\exists B \in \Sigma)(\neg(1 \in v(B))))$$

Eine Aussage A folgt aus einer Prämissenmenge Σ in **LP** genau dann, wenn im Falle, dass alle Prämissen mindestens wahr sind, auch A mindestens wahr ist. Dadurch gilt folgendes Metatheorem bezüglich **LP**:

(LPMT1) **LP** ist eine Erweiterung der Standard-Aussagenlogik.

Begründung: Durch (LPG) können höchstens mehr Aussagen einen ausgezeichneten Wahrheitswert erhalten, alle Standard-Tautologien behalten ihren Wert $\{1\}$ und bleiben also logisch wahr. ■

Veranschaulichen wir dies am Satz vom Ausgeschlossenen Dritten (TND). In Lukasiewiczs dreiwertiger Logik gilt dieser Satz nicht, wie die folgende Wahrheitstafel zeigt:

A	$\vee$	$\neg A$
1	1	0 1
0	1	1 0
u	u	u u

Da "u" kein ausgezeichnete Wahrheitswert ist, gibt es einen Fall, in dem das TND keinen ausgezeichneten Wahrheitswert besitzt, also ein GegenModell, also gilt TND in Lukasiewiczs Logik nicht. In **LP** hingegen:

A	$\vee$	$\neg A$	A
1	1	0	1
0	1	1	0
0,1	0,1	0,1	0,1

Auch in der letzten Zeile gilt $1 \in v(A \vee \neg A)$, d.h. das TND ist in **LP** gültig.

Theoreme in **LP** als einem aussagenlogischen System lassen sich durch die Wahrheitstafel Methode beweisen.

Wir beweisen also (LPT1) wie folgt:

A	B	$\neg A \vee B$	$A \supset B$
1	1	0 1	1
1	0	0 0	0
1	0,1	0 0,1	0,1
0	1	1 1	1
0	0	1 1	1
0	0,1	1 1	1
0,1	1	0,1 1	1
0,1	0	0,1 0,1	0,1
0,1	0,1	0,1 0,1	0,1

Wie in der Standard-Aussagenlogik liefert uns die Wahrheitstafelmethode ein Entscheidungsverfahren für LP-Aussagen. Aus der Existenz eines Entscheidungsverfahrens ergibt sich zugleich die Feststellung:

(LPMT2) **LP** ist vollständig.

Begründung: Allein durch das Auswerten von Wahrheitstafeln können wir alle LP-Wahrheiten erzeugen. ■

Vollständig ist hier genaugenommen nicht der Kalkül **LP**, da wir noch gar keine Beweistheorie behandelt haben, sondern schon bezüglich der wohlgeformten Aussagen von **LP** können wir rein mechanisch feststellen, welche von ihnen logisch wahr sind, welche nicht logisch wahr sind.

Nun haben wir aber mit (LPMT1) gesagt, dass alle Standard-Tautologien auch in **LP** gültig sind. Daraus ergeben sich:

(LPT3) $A \wedge (A \supset B) \supset B$

(LPT4) $(A \supset (A \supset B)) \supset (A \supset B)$

Während die Modus Ponens-Bedingung in gewissem Maße durch (LPT3) als erfüllt betrachtet werden kann, verhindern diese beiden Theoreme das Erfüllen der starken Anti-Trivialitätsbedingung.

Aber droht nicht noch Schlimmeres? Wenn doch alle Standard-Tautologien gelten, dann gilt auch:

(LPT5) $A \wedge \neg A \supset B$

Tatsächlich ist diese Aussage LP-gültig. Heißt dies nun, dass LP überhaupt die schwache Anti-Trivialitätsbedingung nicht erfüllt? Nein, denn (LPMT1) betrifft nur die LP-Gültigkeit. Alle Standard-Tautologien sind auch LP-gültig. Hingegen gelten viele Standard-Folgerungen nicht! Das heißt, dass sich die Folgerungsbeziehung " $\models_{LP}$ " anders verhält als das Konditional " $\supset$ ". Die Folgerungsrelation gilt u.a. bei folgenden Fällen⁷⁵:

$$A \models_{LP} A \vee B \quad A, B \models_{LP} (A \wedge B) \quad A \supset B \models_{LP} (\neg B \supset \neg A)$$

$$\neg \neg A \models_{LP} A \quad A \supset (B \supset C) \models_{LP} B \supset (A \supset C) \quad A \wedge B \models_{LP} B$$

das heißt Doppelte Negation, Konjunktionseinführung und -beseitigung, Disjunktionseinführung, Permutation und Kontraposition gelten auch als Folgerungsregeln. Hingegen gelten in **LP** *nicht*⁷⁶:

$$A \wedge \neg A \models B \quad A \supset B, B \supset C \models A \supset C \quad A, \neg A \vee B \models B$$

$$\hat{A}, A \supset B \models B \quad A \supset B, \neg B \models \neg A \quad A \supset (B \wedge \neg B) \models \neg A$$

das heißt es gelten weder das *ex contradictione quod libet* - was erfreulicherweise heißt, das **LP** die schwache Anti-Trivialitätsbedingung erfüllt - noch - was weit weniger erfreulich ist - Transitivität, Disjunktionsbeseitigung, Modus Ponens(!), Modus Tollens und Negationseinführung als Folgerungsregeln. Ganz allgemein gilt, dass im Falle, dass die Konsequenz und die Prämissen keine Teilaussagen gemeinsam haben, die Folgerung widerlegbar ist⁷⁷. Als Beispiel soll der Disjunktive Syllogismus (die Disjunktionsbeseitigung) widerlegt werden:

Den Fall, dass A nur falsch ist, können wir ausschließen, da dann die Folgerung nach (LPF) trivial wahr ist. A und

⁷⁵ vgl. Priest, "The Logic of Paradox", S.228.

⁷⁶ vgl. ebd.

⁷⁷ vgl. ebd. S.229; insofern weist **LP** ein Merkmal Relevanter Logiken auf, bezogen auf die Folgerungsrelation.

$\neg A \vee B$ müssen bei einem Gegenbeispiel einen ausgezeichneten Wahrheitswert besitzen, wobei B trotzdem nur falsch sein müsste. Dies ist möglich: Wenn A antinomisch ist, dann gilt $1 \in v(A)$ und aufgrund der Wahrheitstafel der Negation

$1 \in v(\neg A)$ sowie aufgrund der Wahrheitstafel der Disjunktion $1 \in v(\neg A \vee B)$ unabhängig davon, welchen Wahrheitswert B hat. Geben wir B also den Wert $\{0\}$! Dann sind beide Prämissen wahr, die vermeintliche Konklusion aber falsch, also besteht nach (LPF) keine Folgerung, der Disjunktive Syllogismus gilt nicht. ■

Obwohl **LP** also die schwache Anti-Trivialitätsbedingung erfüllt, ist dies teuer erkaufte: der Modus Ponens gilt nicht als Folgerungsregel. Wir können mit **LP** also bestimmte Aussagen als logisch wahre Konditionale auszeichnen, wir sind aber nicht in der Lage, unter Hinzuziehung ihrer Antezedenzen mittels Modus Ponens aus ihnen zu folgern! Trotz der Gültigkeit des Modus Ponens-Theorems könnte man daher geneigt sein, **LP** eine Nichterfüllung der Modus Ponens-Bedingung zuzuschreiben.

In **LP** gilt indessen das Deduktionstheorem:

$$(LPMT3) \ A_1 \dots A_n \models_{LP} B \rightarrow A_1 \dots A_{n-1} \models_{LP} (A_n \supset B)$$

Begründung: Wenn der Vordersatz von (LPMT3) wahr ist, muss die Folgerungsbeziehung zwischen $A_1 \dots A_n$ und B bestehen, d.h. jede der Prämissen und B ist mindestens wahr, also ist auch A_n mindestens wahr.

Dann ist aber nach der Wahrheitstafel für das Konditional auch $A_n \supset B$ mindestens wahr und dann gilt nach (LPF) der Hintersatz von (LPMT3). ■

Bezüglich einer Textpassage, die wir für ein vermeintliches Argument mit n Prämissen halten, können wir durch n-maliges Anwenden des Deduktionstheorems ein Konditional bilden. Ist dieses nicht logisch wahr, dann kann nach (LPMT3) auch das zu beurteilende Argument keine Folgerung sein. Wir besitzen so einen *Negativtest für Folgerungen*. Aus der

logischen Wahrheit des so gebildeten Konditionals können wir aber nichts schließen⁷⁸. Wir verfügen über *keinen Positivtest* für Folgerungen.

Als vollständiges System stört uns an **LP** vielleicht die Mühseligkeit der Entscheidung durch Wahrheitstafeln. Wir brauchen ein zeitlich effizienteres Entscheidungsverfahren. Bloesch hat ein solches vorgelegt: eine Tableau-Methode⁷⁹. Eine Tableau-Methode bzw. ein Baum-Kalkül ist ein indirektes Verfahren. Bezüglich einer zu prüfenden Aussage wird angenommen, sie sei falsch (in unserem Fall: nur falsch). Dann wird eine Belegung der sie aufbauenden Teilaussagen gesucht, die dies möglich macht. Dazu wird die logische Komplexität der Aussage schrittweise abgebaut, indem gesagt wird, welche Werte die sie aufbauenden elementaren Aussagen haben müssen, um die Ausgangsbewertung durchzuhalten. Tritt dabei ein Widerspruch auf (d.h. eine unmögliche Bewertung einer Aussage als der Zwang ihr zwei nicht kombinierbare Wahrheitswerte zuzuweisen) dann gibt es die gesuchte Bewertung nicht. Also kann die Ausgangsaussage nicht falsch gemacht werden. Also muss sie logisch wahr sein. Tritt kein Widerspruch auf, ist die Ausgangsaussage nicht logisch wahr. Das Verfahren läßt sich auch auf die Überprüfung von Folgerungsbeziehungen anwenden, indem man den Prämissen den Wert "mindestens wahr" zuschreibt und der Konklusion den Wert "bestimmt nicht wahr" und dann beginnt auszuwerten.

Die Baum-Regeln ergeben sich aus den Wahrheitstafeln. Es handelt sich um ein semantisch motiviertes Verfahren. Bloesch betrachtet **LP** nun als 4-wertiges System mit den Werten: "bestimmt nicht wahr" (d.h. nur falsch: "T"), "mindestens wahr" ("T"), "bestimmt nicht falsch" (d.h. nur wahr: "F") und "mindestens falsch" ("F"). Als unvereinbar gelten die Wahrheitswerte nach folgender Tafel:

⁷⁸ Dies wäre der Trugschluß der Bejahung des Hintersatzes.

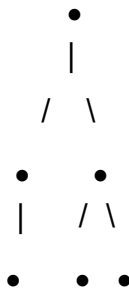
⁷⁹ vgl. Bloesch, "A Tableau Style Proof System for Two Paraconsistent Logics".

$T(A)$ (mindestens wahr)	$T'(A)$ (gar nicht wahr = nur falsch)
$T'(A)$ (gar nicht wahr = nur falsch)	$F'(A)$ (gar nicht falsch = nur wahr)
$F(A)$ (mindestens falsch)	$F'(A)$ (gar nicht falsch = nur wahr)

Es ergeben sich beispielsweise folgende Regeln der Negation:

<u>$T(\neg A)$</u>	<u>$F(\neg A)$</u>	<u>$T'(\neg A)F'(\neg A)$</u>	
$F(A)$	$T(A)$	$F'(A)$	$T'(A)$

Der Strich besagt, dass unterhalb einer Aussage des Typs, der oberhalb des Strichs steht, eine Aussage angehängt werden muss des Typs, der unterhalb des Strichs steht. Bei diesen taucht die Negation nicht mehr auf. Logische Komplexität wurde also abgebaut. Findet sich unter des Strichs auch noch eine senkrechte Trennwand, so heißt dies, dass unterhalb einer bestimmten Aussage zwei Äste eine Alternative bilden. Es ergeben sich so Gebilde wie folgendes:



wobei die Punkte Aussagen symbolisieren.

Jetzt kann man definieren:

(LPB1) Ein Ast eines Baumes heißt erfüllbar, wenn alle Aussagen, die von ihm aus von unten nach oben erreichbar sind, einen ausgezeichneten Wahrheitswert erhalten können in einer Bewertung v .

(LPB1') Ein Ast eines Baumes heißt erfüllbar, wenn keine Aussage, die

von ihm aus von unten nach oben erreichbar ist, mit widersprechenden Wahrheitswerten belegt ist.

Diese Formulierung bezieht sich auf die Tabelle der sich ausschließenden Wahrheitswerte.

(LPB2) Ein Ast, der nicht erfüllbar ist, ist geschlossen.

(LPB3) Ein Baum ist erfüllbar, wenn mindestens ein Ast erfüllbar ist.

(LPB4) Ein Baum ist geschlossen, wenn alle seine Äste geschlossen sind.

(LPB5) Eine Aussage A ist logisch wahr in LP bzw. eine Folgerung $\Sigma \models_{LP} B$ gültig, wenn ihr Baum geschlossen ist.

Entscheidend ist die letzte Bestimmung. Bloesch zeigt dann

(LPMT4) (i) $\models_{LP} A \leftrightarrow$ Der Baum mit dem Ausgangspunkt $T'(A)$ ist geschlossen.

(ii) $A_1 \dots A_n \models_{LP} B \leftrightarrow$ Der Baum mit dem Ausgangspunkt $T(A_1) \dots T(A_n), T'(B)$ ist geschlossen

Damit liefert das Baumverfahren ein Entscheidungsverfahren für LP-Gültigkeit und einen weiteren Vollständigkeitsbeweis.

Für die Junktoren gelten folgende Baum-Regeln:

<u>$T(\neg A)$</u>	<u>$F(\neg A)$</u>	<u>$T'(\neg A)F'(\neg A)$</u>	
$F(A)$	$T(A)$	$F'(A)$	$T'(A)$

<u>$T(A \wedge B)$</u>	<u>$F(A \wedge B)$</u>	<u>$T'(A \wedge B)$</u>	<u>$F'(A \wedge B)$</u>
TA	$FA \mid FB$	$T'A \mid T'B$	$F'A$
TB			$F'B$

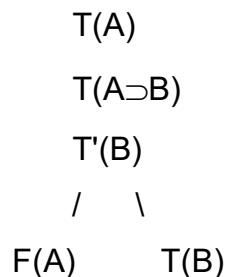
<u>$T(A \vee B)$</u>	<u>$F(A \vee B)$</u>	<u>$T'(A \vee B)$</u>	<u>$F'(A \vee B)$</u>
$TA \mid TB$	FA	$T'A$	$F'A \mid F'B$
	FB	$T'B$	

<u>T(A$\supset$B)</u>	<u>F(A$\supset$B)</u>	<u>T'(A$\supset$B)</u>	<u>F'(A$\supset$B)</u>
FA TB	TA	F'A	T'A F'B
	FB	T'B	

<u>T(A$\equiv$B)</u>	<u>F(A$\equiv$B)</u>	<u>T'(A$\equiv$B)</u>	<u>F'(A$\equiv$B)</u>
FA TA	TA FA	F'A T'A	T'A F'A
FB TB	FB TB	T'A F'B	T'B F'B

Zur besseren Lesbarkeit werden auch Klammern verwendet. Betrachten wir nun zwei Beispiele:

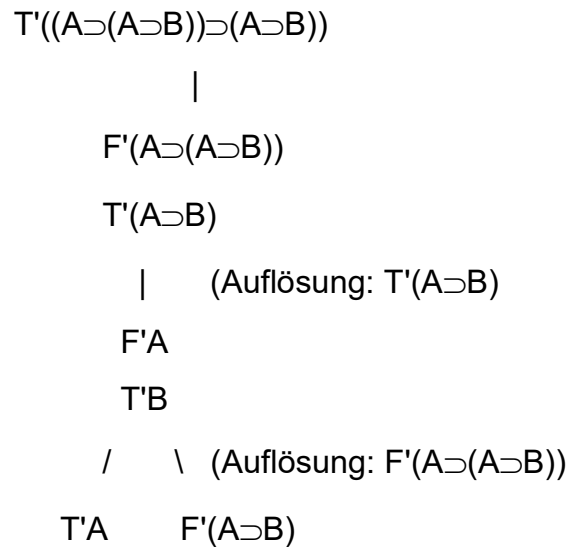
Zunächst zeigt der Baum-Kalkül, dass Modus Ponens als Folgerung nicht gilt:



Der rechte Ast dieses Baums ist nach (LPB1') geschlossen, da dort sowohl T(B) als auch T'(B) vorkommt. Der linke Ast ist hingegen offen, da sich "mindestens wahr" und "mindestens falsch" bezüglich A nicht widersprechen. Also ist nach (LPB4) der Baum nicht geschlossen. Also ist nach (LPB5) die Folgerung nicht gültig. ■

Bei **LP** handelt es sich also um einen Kalkül, in dem Modus Ponens *als Regel* nicht vorliegt. Damit wird eine der Bedingungen, welche die Formen des Curry Paradox auszeichnen negiert. Damit lassen sich die Curry Paradoxien nicht ableiten. Damit ist **LP** nicht trivial bezüglich der naiven Semantik!

Gültig hingegen ist die Absorption:



Der linke Ast ist geschlossen, da sich auf ihm $T'A$ und $F'A$ befinden. Der rechte Ast ist geschlossen, da sich auf ihm $T'(A \supset B)$ und $F'(A \supset B)$ befinden. Also ist der Baum geschlossen. Also ist Absorption gültig. ■

Die Gültigkeit der Absorption - und analog des Modus Ponens-Theorems kann aber in **LP** keinen Schaden anrichten, da ja eine andere Bedingung der Curry Paradoxien fehlt.

Diese Einfachheit von **LP** und des zugehörigen Baumverfahrens kann der Anlass sein, **LP** prädikatenlogisch zu **LPQ** erweitern zu wollen. Nehmen wir dazu eine einfache Quantorenlogik mit einem festen Bereich und ohne Kennzeichnungen⁸⁰.

Die Semantik der Quantorensprache, die ich hier ausführe, ist nicht modal. Der Bereich D der Individuen muss also als der gleichbleibende Bereich der existenten oder der möglichen Gegenstände verstanden werden. Diese basale Quantorensprache enthält auch kein Existenzprädikat und keine Kennzeichnungen als Individuenausdrücke. Diese und entsprechende Quantoren können jedoch auf die übliche Weise hinzuge-

⁸⁰ Bzw. müßten Aussagen, in denen Kennzeichnungen vorkommen, nur in ihre kennzeichnungsfreie Quantorenform gebracht werden, um nach den folgenden Regeln behandelt werden zu können.

fügt werden⁸¹. Das Vorgehen einer solchen Erweiterung bzw. Modifikation unterscheidet sich dann nicht vom Übergang von der Standard-Logik zur modalen oder Freien Logik. Parakonsistenz spielt hier keine Rolle. Das Abblocken von Trivialität (das parakonsistente Element in einer Logik) ist ein rein aussagenlogisches Problem. Deshalb beschränke ich mich in dieser Arbeit, in der es nicht um Existenz oder Kennzeichnungen geht, auf die elementare Quantorenlogik (evtl. - wie bei **SKP** - mit Identität).

In einer prädikatenlogischen Sprache verfügen wir über neben den Ausdrücken, die in **LP** vorkommen, über Prädikatoren " $F^1_1()$ "...

" $F^n_1()$ "..." $F^m_n()$ " (wobei anstelle der Indizierung für bestimmte Prädikatoren auch " $G()$ ", " $H()$ " usw. geschrieben werden kann), die eine bestimmte Anzahl von - gerade oben indizierten - Leerstellen besitzen, in die Individuenausdrücke eingesetzt werden, das Identitätssymbol "=", Individuenkonstanten " a_1 "..." a_n " (bzw. " b ", " c " usw.), Individuenvariablen " x_1 "... (bzw. " y ", " z " usw.) sowie die Quantoren " $\forall$ " und " $\exists$ ". Hinzufügen können wir den Kennzeichnungsoperator " ιx " sowie das Existenzprädikat " $E()$ ", wenn wir auf diese Weise weitere Quantoren (z.B. für existente im Unterschied zu möglichen Objekten) einführen wollen, was hier zunächst nicht geschehen soll.

Ein **LPQ-Modell** m ist ein Tupel $\langle D, v, G \rangle$ mit einem Individuenbereich D , einer Interpretationsfunktion v und der Menge der Variablenbelegungen G . Die Modelle m_i in der Menge der Modelle M unterscheiden sich durch den Individuenbereich oder die Interpretation.

Es gelten neben den Wahrheitstafeln für aussagenlogisch komplexe Aussagen folgende semantische Regeln:

(LPQS1) Sei α eine Individuenkonstante: $v(\alpha) = d \in D$

Jeder Individuenkonstante wird ein Element des Gegenstandsbereiches zugeordnet.

⁸¹ vgl. z.B. Kutschera, *Einführung in die intensionale Semantik*.

(LPQS2) Sei γ ein n -stelliger Prädikator: $v(\gamma) = \langle P^+, P^- \rangle$ wobei P^+, P^- Teilmengen von D^n sind.

Jedem n -stelligen Prädikator wird als Extension ein Tupel zugewiesen, das besteht aus einem Positivbereich aus D^n (dem n -fachen Cartesischen Produkt von D , d.h. der Menge der n -Tupel aus D), der Gegenstände, auf die $P()$ uneingeschränkt zutrifft, sowie einem Negativbereich der Gegenstände, auf die $P()$ uneingeschränkt nicht zutrifft. Gegenstände, die sich bezüglich $P()$ antinomisch verhalten, sind in beiden Bereichen.

Aus (LPQS1) und (LPQS2) ergibt sich die wesentliche Wahrheitsbedingung für Elementaraussagen:

(LPQS3) Sei $P^n()$ ein n -stelliger Prädikator und α^n ein n -Tupel von

Individuenkonstanten, dann bildet $v \langle P^n(), \alpha^n \rangle$ auf $\{\{1\}, \{0\}, \{0, 1\}\}$ ab,

wobei $d^n = v(\alpha^n)$, je nachdem ob $d^n \in P^+$, $d^n \in P^-$, ($d^n \in P^+$ und $d^n \in P^-$).

(bzw. $\{F, T, T \& F\}$ in Bloesch's-Diktion).

Das heißt eine elementare Aussage hat einen bestimmten semantischen Wert genau dann, wenn das Tupel von Gegenständen, auf die sich, die in ihr vorkommenden Individuenkonstanten beziehen, in den Bereich der Interpretation des generellen Terms fällt.

Aussagenbuchstaben können auch von v interpretiert werden, indem sie auf $\{\{1\}, \{0\}, \{0, 1\}\}$ bewertet werden.

(LPQS4) Für $g \in G$ ist $g(x_i) = d \in D$.

Eine Variablenbelegung ordnet einer Variablen irgendein Objekt aus dem Individuenbereich zu. Ist $P(x_1 \dots x_n)$ eine offene Formel (d.h. in dem Prädikator kommen Variablen vor, die nicht im Skopus eines Quantors stehen), dann bezieht sich in diesem Fall die Interpretationsfunktion auf $\langle P(x_1 \dots x_n), \langle g(x_1) \dots g(x_n) \rangle \rangle$, d.h. sie bezieht sich auf die Variablenbelegungen der frei vorkommenden Variablen. Wenn zwei Variablenbelegungen übereinstimmen bis höchstens auf die Belegung einer Variablen " x " ("bis auf die Stelle x "), dann hängt ihre unterschiedliche Bewertung von Aussagen nur daran, dass sie " x " verschieden bewertet haben.

Über die wahrheitsfunktionalen Ausdrücke hinaus müssen wir nun semantische Regeln für die Quantoren angeben:

" $v(\alpha, g)$ " bedeutet, dass v die Aussage bzw. Aussageformel relativ zur Variablenbelegung g bewertet.

(LPQS5) Sei α eine Aussage der Form $(\forall x)P(x)$, dann gilt:

- (i) $v(\alpha, g) = \{1\} \leftrightarrow (\forall g' \in G)(g' \text{ unterscheidet sich von } g \text{ höchstens an der Stelle } x \text{ und } g'(x) \in P^+)$.
- (ii) $v(\alpha, g) = \{0\} \leftrightarrow (\exists g' \in G)(g' \text{ unterscheidet sich von } g \text{ höchstens an der Stelle } x \text{ und } g'(x) \in P^-)$
- (iii) $v(\alpha, g) = \{0, 1\} \leftrightarrow (\forall g' \in G)(g' \text{ unterscheidet sich von } g \text{ höchstens an der Stelle } x \text{ und } g'(x) \in P^+ \text{ und } (\exists g'' \in G)(g'' \text{ unterscheidet sich von } g \text{ höchstens an der Stelle } x \text{ und } g''(x) \in P^-)$

Klausel (i) besagt, dass eine Allaussage wahr ist, wenn für die Variable, die allquantifiziert ist, Beliebiges eingesetzt werden kann, sie also von jedem beliebigen Gegenstand, also von allen Gegenständen gilt, was sich darin ausdrückt, dass alle Variablenbelegungen, die sich nur an der Stelle für "x" unterscheiden, und die, da es sich um alle Variablenbelegungen handelt, zusammen alle Gegenstände für "x" einsetzen, diese Aussage wahr machen. Klausel (ii) verweist auf die Existenz eines Gegenbeispiels zu einer Allaussage. Eine Allaussage ist antinomisch, wenn sie bezüglich keines Gegenstands allein falsch, aber bezüglich mindestens eines Gegenstandes antinomisch ist.

Da für den Existenzquantor das Theoremschema⁸² gilt:

$$(LPQT1) (\exists x)P(x) \equiv \neg(\forall x)\neg P(x)$$

ergeben sich entsprechende Klauseln.

⁸² Da wir uns schon tief im Gebiet der Parakonsistenten Logik befinden, wo es keine Heterogenität von Objekt- und Metasprache gibt, werde ich einige Beweise gemäß dem Standard-Vorgehen für Schemata durchführen, wie dies bis jetzt unter Verwendung von "A" und "B" usw. immer geschehen ist, aber zugleich dazu übergehen, Beweise für einzelne Aussagen und Folgerungszusammenhänge der Objektsprache mittels homophoner Übersetzungen zu führen.

Eine Aussage bzw. Folgerung ist LPQ-gültig nach der Definition:

$$(LPQG) \models_{LPQ} \leftrightarrow (\forall m_i \in M)(1 \in v_i(A))$$

Eine Aussage ist LPQ-gültig, wenn sie in allen Modellen gemäß deren Interpretationsfunktion mindestens wahr ist.

$$(LPQF) \Sigma \models_{LPQ} A \leftrightarrow (\forall m_i \in M)(\exists B \in \Sigma)(v_i(B) = \{0\} \vee 1 \in v_i(A))$$

Eine Folgerung ist LPQ-gültig, wenn für jedes Modell gilt, dass es, wenn es entweder eine der Prämissen ausschließlich falsch macht oder die Konklusion mindestens wahr macht.

Jetzt soll ein Baum-Kalkül **LPQB** definiert werden. **LPQB** setzt sich zusammen aus den Baumregeln für **LP** und den folgenden Regeln für die Quantoren:

$$\forall \quad \frac{\underline{T(\forall x)P(x)}}{T(P(a))^+} \quad \frac{\underline{F(\forall x)P(x)}}{F(P(a))^*} \quad \frac{\underline{T'(\forall x)P(x)}}{T'(P(a))^*} \quad \frac{\underline{F'(\forall x)P(x)}}{F'(P(a))^+}$$

* "a" muss eine Individuenkonstante sein, die im Baum bis dahin nicht vorkam.

+ diese Regel kann innerhalb eines Baumes mehrmals zur Anwendung kommen.

$$\exists \quad \frac{\underline{T(\exists x)P(x)}}{T(P(a))^*} \quad \frac{\underline{F(\exists x)P(x)}}{F(P(a))^+} \quad \frac{\underline{T'(\exists x)P(x)}}{T'(P(a))^+} \quad \frac{\underline{F'(\exists x)P(x)}}{F'(P(a))^*}$$

* "a" muss eine Individuenkonstante sein, die im Baum bis dahin nicht vorkam.

+ diese Regel kann innerhalb eines Baumes mehrmals zur Anwendung kommen.

Die Einschränkungen mit "*" entsprechen den üblichen Einschränkungen der existentiellen Spezialisierung. Aufgrund der Quantorendualität (LPQT1) treten sie bei beiden Gruppen von Regeln auf. Die mehrmalige

Anwendung von Regeln ergibt sich aus dem allgemeinen Charakter der betreffenden Aussage, die auch auf im Baum - durch Existenzspezialisierungen - neu auftretende Individuenkonstanten anzuwenden sind.

Die Definitionen (LPB1) bis (LPB5) können entsprechend zu (LPQB1) bis (LPQB5) übernommen werden.

Betrachten wir Beispiele von quantorenlogischen Baumableitungen:

(LPQB1) $(\forall x)F(x), (\forall x)G(x) \models_{LPQ} (\forall x)(F(x) \wedge G(x))$

Beweis:

- | | | |
|----|-----------------------------------|----------------------------|
| 1. | $T(\forall x)F(x)$ | |
| 2. | $T(\forall x)G(x)$ | |
| 3. | $T'(\forall x)(F(x) \wedge G(x))$ | |
| | | (Auflösung von 3) |
| 4. | $T'(F(a) \wedge G(a))$ | ("a" ist neu) |
| | / \ | |
| | $T'F(a)$ | $T'G(a)$ (Auflösung von 4) |
| | | (Auflösung 1, 2) |
| | $T F(a)$ | $T G(a)$ |

Beide Äste des Baumes sind geschlossen. Bei der Auflösung der beiden mindestens wahren Allaussagen wird auf die im Baum vorkommende Individuenkonstante "a" zurückgegriffen, wodurch unmittelbar die Schließung erfolgt. Also ist der Baum geschlossen (LPQB4), also ist nach (LPQB5) die zu prüfende Folgerung gültig.

(LPQB2) $(\forall x)(F(x) \supset G(x)) \models_{LPQ} (\forall x)F(x) \supset (\forall x)G(x)$

Beweis:

1. $T((\forall x)(F(x) \supset G(x)))$
2. $T'((\forall x)F(x) \supset (\forall x)G(x))$
 $\quad \quad \quad | \quad \quad \quad$ (Auflösung 2)
3. $F'(\forall x)F(x)$
4. $T'(\forall x)G(x)$
 $\quad \quad \quad | \quad \quad \quad$ (Auflösung 4)
5. $T'G(a)$
 $\quad \quad \quad | \quad \quad \quad$ (Auflösung 3)
6. $F'F(a)$
 $\quad \quad \quad | \quad \quad \quad$ (Auflösung 1)
7. $T(F(a) \supset G(a))$
 $\quad \quad \quad / \quad \quad \quad \backslash$ (Auflösung 7)
8. $F(F(a)) \quad T(G(a))$

Auch hier sind beide Äste des Baums geschlossen, da Zeile 8 in Widerspruch gerät mit Zeile 6 (links) und Zeile 5 (rechts), also ist der Baum geschlossen, also ist die Folgerung gültig.

Ebenfalls gültig sind:

(LPQB3) $(\forall x)(F(x) \supset G(x)) \models_{LPQ} (\forall x)(\neg G(x) \supset \neg F(x))$

(LPQB4) $(\forall x)(F(x) \supset G(x)) \models_{LPQ} (\exists x)F(x) \supset (\exists y)G(y)$

Um zugleich eine einfachere Darstellung einzuführen, die sich an herkömmlichen Ableitungen orientiert, seien noch bewiesen⁸³:

⁸³ Dieser Beweis findet in der Objektsprache statt.

(LPQB5) $(\forall x)F(x) \supset (\exists x)F(x)$

- Beweis:
- | | | |
|----|---|-----------------|
| 1. | $T'((\forall x)F(x) \supset (\exists x)F(x))$ | |
| 2. | $F'(\forall x)F(x)$ | $T' \supset, 1$ |
| 3. | $T'(\exists x)F(x)$ | $T' \supset, 1$ |
| 4. | $T'F(a)$ | $T' \exists, 3$ |
| 5. | $F'F(a)$ | $F' \forall, 2$ |
| 6. | # | |

"#" stellt hier fest, dass eine unmögliche Bewertung auftritt (im Beispiel zwischen (4) und (5)). Rechts neben den Zeilen werden die angewendeten Regeln angegeben, und die Zeilen, auf die sie angewendet wurden.

(LPQB6) $(\forall x)(P(x) \supset A) \models_{LPQ} (\exists x)P(x) \supset A$

wobei "x" in "A" nicht frei vorkommt.

- Beweis⁸⁴:
- | | | |
|-----|---------------------------------|--|
| 1. | $T(\forall x)(P(x) \supset A)$ | |
| 2. | $T'((\exists x)P(x) \supset A)$ | |
| 3. | $F'(\exists x)P(x)$ | $T' \supset, 2$ |
| 4. | $T'A$ | $T' \supset, 2$ |
| 5. | $F'P(\acute{a})$ | $F' \exists, 3$ $\acute{a}$ ist neu |
| 6. | $T(P(\acute{a}) \supset A)$ | $T \forall, 1$ |
| 7a. | $F(P(\acute{a}))$ | $T \supset, 6$ 7b. TA $T \supset, 6$ |
| 8a. | # | 5, 7a. 8b. # 4, 7a. ■ |

Entsprechend:

(LPQB7) $(\exists x)(F(x) \wedge G(x)) \models_{LPQ} (\exists x)F(x) \wedge (\exists x)G(x)$

Nicht gültig sind hingegen die quantorenlogischen Entsprechungen zu den ungültigen aussagenlogischen Folgerungen, d.h. nicht gültig sind u.a.

- (i) $(\forall x)P(x), (\forall x)(P(x) \supset Q(x)) \models (\forall x)Q(x)$
 (ii) $(\forall x)\neg Q(x), (\forall x)(P(x) \supset Q(x)) \models (\forall x)\neg P(x)$

⁸⁴ Dieser Beweis hingegen verwendet die Schemata (z.B. auch $\acute{a}$).

(iii) $(\forall x)(P(x) \supset Q(x)), (\forall x)(P(x) \supset R(x)) \models (\forall x)(P(x) \supset R(x))$

(i) entspricht dem Modus Ponens, (ii) dem Modus Tollens, (iii) der Transitivität. Die Ungültigkeit zeigt sich daran, dass ein Baum sich nicht schließt;

ad (i):

- | | |
|---|---|
| 1. $T(\forall x)(P(x) \supset Q(x))$ | |
| 2. $T(\forall x)P(x)$ | |
| 3. $T'(\forall x)Q(x)$ | |
| 4. $T'Q(\acute{a})$ | $T'\forall, 3 \text{ } \acute{a} \text{ ist neu}$ |
| 5. $T P(\acute{a})$ | $T\forall, 2$ |
| 6. $T(P(\acute{a}) \supset Q(\acute{a}))$ | $T\forall, 1$ |
| 7a. $F(P(\acute{a})) \quad T\supset, 6$ | 7b. $T(Q(\acute{a})) \quad T\supset, 6$ |
| | 8b. $\# \quad 4, 7b.$ |

Während der rechte Ast geschlossen ist, bleibt der linke Ast offen, da es sich bei $F(P(\acute{a}))$ und $T(P(\acute{a}))$ nicht um eine widersprüchliche Bewertung handelt: $P(\acute{a})$ kann zugleich mindestens falsch und mindestens wahr sein, nämlich wenn $P(a)$ antinomisch ist.

ad (ii):

- | | | |
|---|---|---|
| Beweis: | 1. $T(\forall x)(P(x) \supset Q(x))$ | |
| | 2. $T(\forall x)\neg Q(x)$ | |
| | 3. $T'(\forall x)\neg P(x)$ | |
| | 4. $T'\neg P(\acute{a})$ | $T'\forall, 3 \text{ } \acute{a} \text{ ist neu}$ |
| | 5. $T \neg Q(\acute{a})$ | $T\forall, 2$ |
| | 6. $T(P(\acute{a}) \supset Q(\acute{a}))$ | $T\forall, 1$ |
| | 7. $F'P(\acute{a})$ | $T'\neg, 4$ |
| | 8. $F Q(\acute{a})$ | $T\neg, 5$ |
| 9a. $F(P(\acute{a})) \quad T\supset, 6$ | 9b. $F(Q(\acute{a})) \quad T\supset, 6$ | |
| 10a. $\# \quad 7, 9a$ | | |

Während sich der linke Ast schließt, bleibt der rechte Ast offen, so dass der quantorenlogische Modus Tollens in **LPQB** nicht gültig ist.
(Entsprechend für (iii).)

Die Korrektheit von **LPQB** ergibt sich daraus, dass im Falle, dass ein Baum für eine Aussage α geschlossen ist, dies besagt, dass α nicht mit einer Bewertung v versehen werden kann, so dass sich aus der Bewertung der α aufbauenden elementaren Aussagen und den Wahrheitsbedingungen der logischen Ausdrücke ergibt, dass α nur falsch ist. α muss also mindestens wahr sein (entweder bloß wahr oder antinomisch). Da aber sowohl $\{1\}$ als auch $\{0,1\}$ ausgezeichnete Wahrheitswerte sind, besitzt α in jedem Fall einen ausgezeichneten Wahrheitswert, d.h. es gibt keine Bewertung, die α einen nicht-ausgezeichneten Wahrheitswert gibt, d.h. nach Definition der **LPQ**-Gültigkeit α ist **LPQ**-gültig. Als Metatheorem:

(LPQBMT1) (i) Der **LPQ**-Baum von $T'A$ ist geschlossen $\rightarrow \models_{\text{LPQ}} A$

und entsprechend für Folgerungen:

(LPQBMT1) (ii) Der **LPQ**-Baum mit dem Ausgangspunkt $T(A_1) \dots T(A_n)$, $T'(B)$ ist geschlossen $\rightarrow A_1 \dots A_n \models_{\text{LPQ}} B$

Wie sieht es mit der deduktiven Vollständigkeit aus? **LPQB** wurde aus **LPB** gewonnen, indem letzterer mit quantorenlogischen Regeln versehen wurde, wobei die verwendete Quantorensemantik der Standard-Quantorensemantik entsprach. Das heißt aber: Der übliche metalogische Beweis, dass ein vollständiger aussagenlogischer Baumkalkül durch eine Erweiterung um die Standard-Quantoren zu einem vollständigen prädikatenlogischen Baumkalkül wird⁸⁵, muss sich - mit einigen Zusatzklauseln für den Fall antinomischer Bewertung - auf **LPQB**

⁸⁵ vgl. z.B. Ben-Ari, *Mathematical Logic for Computer Science*, 112-117.

übertragen lassen. Da nun Bloesch gezeigt hat, dass **LPB** vollständig ist, ergibt sich:

(LPQBMT2) **LPQB** ist deduktiv vollständig:

$\models_{\text{LPQA}} A \rightarrow$ Der LPQ-Baum von T'A ist geschlossen.

LPQB ist also ein adäquater Kalkül. Außerdem erfüllt **LPQB** nicht nur die extensionale Adäquatheitsbedingung, sondern auch zugleich die schwache und die starke Anti-Trivialitätsbedingung. Was können wir mehr wollen? Nun, **LPQB** erfüllt die beiden Anti-Trivialitätsbedingungen *nur* deshalb, weil die Modus Ponens-Bedingung eigentlich - selbst angesichts der Gültigkeit des Modus Ponens-Theorems - nicht erfüllt ist. **LPQB** besitzt einen konditionalen Junktor, doch wir können vom Wahrsein seines Antezedens nicht auf das Wahrsein seines Konsequenz folgern, es sei denn wir wissen, dass der Kontext nicht antinomisch ist. Für nicht-antinomische Kontexte bräuchten wir aber **LPQB** sowieso nicht. Also haben wir einen konditionalen Junktor, der Folgerung nicht zulässt. Und das ist ein schwerer Anschlag auf unser intuitives Verständnis eines überhaupt akzeptablen "wenn..., dann ---". Die Frage ist, ob uns dieser Preis zu hoch ist für eine nicht-triviale Logik der naiven Semantik⁸⁶.

Besser noch wären Kalküle, die zusätzlich zu den Vorzügen von **LPQB** auch noch die Modus Ponens-Bedingung erfüllen. Sie wären ausgezeichnet, wenn sie die Bedingung der minimalen Beschädigung besser erfüllen als **LPQB**.

⁸⁶ Während Priest anderenorts die Modus Ponens-Bedingung im Sinne des gerade Gesagten geradezu zum definieren einer überhaupt akzeptablen Logik machen möchte, ist der einzige Kalkül, der in (LT) als parakonsistente Logik vorgestellt wird LP (LT:188)! An dieser Stelle behauptet Priest, dies sei die Logik aus (IC:Kap.5), was so schlichtweg falsch ist, da dort *kein* konditionaler Junktor verwendet wird.

Ein parakonsistenter Sequenzen-Kalkül

Als zweiten Kalkül betrachten wir nun einen Sequenzenkalkül. In Nicht-Sequenzenkalkülen geht es um den Beweis von logischen wahren Aussagen (Theoremen) aus anderen logisch wahren Aussagen (insbesondere in axiomatischen Systemen) oder um den Beweis einer Folgerungsbeziehung zwischen einer Prämissenmenge und einer Konklusion (insbesondere in Systemen des Natürlichen Schließens). Den letzteren ähnelt ein Sequenzenkalkül⁸⁷. Sequenzenkalküle verzichten - im Unterschied zu Kalkülen des Natürlichen Schließens - auf Annahmeformeln und ersetzen eine entsprechende Herleitung durch die Sequenz

$$A_1 \dots A_n \Vdash B$$

Es geht darum, ob und wie, wenn von einer Prämissenmenge eine Aussage bewahrheitet wird, die Kombination zweier Prämissenmenge zur Bewahrheitung weiterer Aussagen führt. Die Zeile, in der gesagt wird, dass von einer Prämissenmenge Σ eine Aussage bewährt wird, nennt man eine "Sequenz". Die Regeln eines Sequenzenkalküls beziehen sich darauf, wie man vom Vorliegen mehrerer solcher Sequenzen zu neuen Sequenzen kommt. Einen Spezialfall der Sequenzen bilden die Axiome, da sie aus einer leeren Prämissenmenge folgen. Zu erinnern ist daran, dass die Beziehung, die zwischen einer Prämissenmenge und einer Konklusion aus ihr besteht, die Beziehung der Wahrheitsvererbung ist, wäh-

⁸⁷ vgl. Gentzen, "Untersuchungen über das logische Schließen". Bei Gentzen besitzen die Sequenzenkalküle Einführungsregeln für die logischen Ausdrücke jeweils links und rechts von " $\Vdash$ ", aber keine Beseitigungsregeln. Gentzens Regeln der "Verdünnung", die im folgenden aber gerade nicht vorkommt, kann eine Quelle der Irrelevanz sein, wobei diese Irrelevanz durch eine Restriktion der Beweisführung vermeidbar ist (vgl. Tennant, "Perfect Validity, Entailment and Paraconsistency"), wobei Tennat allerdings auf einen konditionalen Junktor verzichtet und die Sequenz $A \Vdash B$ wie das Enthaltensein behandelt; allerdings fände dann alle Rede über Enthaltensein in der entsprechenden Metasprache statt, es gäbe keine Abtrennung, Transitivität des Enthaltenseins - usw. Hier kommen sowohl " $\Vdash$ " als auch " $\Rightarrow$ " in einem Kalkül vor.

rend " $\Rightarrow$ " das *stärkere* Enthaltensein ausdrückt. (Weswegen auch das Deduktionstheorem im Gegensatz zu **LP** nicht gilt.)

In der Syntax stimmt der Sequenzenkalkül mit der Prädikatenlogik **LPQ** aus dem vorherigen Abschnitt überein. Hinzukommt das " $\Vdash$ " Zeichen, welches besagt, dass sich die Wahrheit der links von diesem Zeichen stehenden Aussagen (den Annahmen oder Prämissen einer Sequenz) auf die Aussage überträgt, die rechts von diesem Zeichen steht, d.h. bei Aufführen aller Bedingungen⁸⁸:

$$(S \Vdash) \text{ (i) } (\gamma_1 \dots \gamma_n \Vdash \alpha) \rightarrow \neg(\exists v \in V)(1 \in v(\gamma_1) \wedge \dots \wedge 1 \in v(\gamma_n) \wedge v(\alpha) = \{0\})$$

$$\text{ (ii) } \neg(\exists v \in V)(1 \in v(\gamma_1) \wedge \dots \wedge 1 \in v(\gamma_n) \wedge v(\alpha) = \{0\}) \rightarrow (\gamma_1 \dots \gamma_n \Vdash \alpha)$$

$$\text{ (iii) } (\exists v \in V)(1 \in v(\gamma_1) \wedge \dots \wedge 1 \in v(\gamma_n) \wedge v(\alpha) = \{0\}) \rightarrow \neg(\gamma_1 \dots \gamma_n \Vdash \alpha)$$

$$\text{ (iv) } \neg(\gamma_1 \dots \gamma_n \Vdash \alpha) \rightarrow (\exists v \in V)(1 \in v(\gamma_1) \wedge \dots \wedge 1 \in v(\gamma_n) \wedge v(\alpha) = \{0\})$$

Die Wahrheitsvererbung tritt nicht auf, wenn bei Wahrheit der linken Seite von " $\Vdash$ " die rechte Seite *nur* falsch wird. Ist die rechte Seite antinomisch, so ist sie *auch* wahr, d.h. die Wahrheitsvererbung liegt vor. Da die rechte Seite entweder "nur falsch" oder nicht "nur falsch" ist kann eine $\Vdash$ -Aussage *niemals* antinomisch sein: Wenn eine Aussage bezüglich Wahrheitsvererbung wahr ist, dann kann bei Wahrheit der linken Seite von " $\Vdash$ " nicht $v(A) = \{0\}$ sein, so dass dann dieselbe Aussage nicht zugleich falsch sein kann, da nun schon $1 \in v(A)$ gelten muss, also Wahrheit auf der rechten Seite vorliegt.

Ein Relevanter Sequenzenkalkül⁸⁹

Die Axiome der Form " $\Vdash \alpha$ " behaupten von einer Aussage α , dass sie ausgehend von einer leeren Prämissenmenge behauptet werden kann

⁸⁸ Warum statt einer Äquivalenz hier alle Bedingungen aufzuführen sind, wird in 4.3 klar.

⁸⁹ Der aussagenlogische Teil des Kalküls stammt aus: Priest, "Sense, Entailment and Modus Ponens", S.426f.

(also logisch wahr ist). Die Regeln beziehen sich darauf, wie aus vorliegenden Sequenzen neue gebildet werden können, wobei durch das Zusammenfassen (Vereinigen) der Prämissen neue Konsequenzen abgeleitet werden können. " Π ", " Σ " stehen für solche Prämissenmengen, sie können leer oder identisch sein.

(Die Erläuterungen zu den einzelnen aussagenlogischen Regeln und Axiomen finden sich schon im Vorhergehenden.)

$$(D1) \quad A \leftrightarrow B := (A \Rightarrow B) \wedge (B \Rightarrow A)$$

$$(D2) \quad A \vee B := \neg(\neg A \wedge \neg B)$$

Da man gemäß den LP-Wahrheitstafeln für " $\neg$ ", " $\wedge$ " und " $\vee$ ", die auch für **SKP** gelten sollen, die Disjunktion durch die Konjunktion definieren kann oder umgekehrt, könnten die folgenden Axiome auch mit nur einem dieser Junktoren formuliert werden.

$$(A1) \quad \{A\} \Vdash A$$

$$(A2) \quad \Vdash A \wedge B \Rightarrow A \quad \& \quad \Vdash A \wedge B \Rightarrow B$$

$$(A3) \quad \Vdash A \Rightarrow A \vee B \quad \& \quad \Vdash B \Rightarrow A \vee B$$

$$(A4) \quad \Vdash A \Rightarrow \neg\neg A$$

$$(A5) \quad \Vdash \neg\neg A \Rightarrow A$$

$$(A6) \quad \Vdash A \wedge (B \vee C) \Rightarrow (A \wedge B) \vee (A \wedge C)$$

$$(A7) \quad \Vdash A \vee \neg A$$

$$(A8) \quad \Vdash (\forall x)x=x$$

Alle Gegenstände sind mit sich identisch.

$$(R1) \quad \Sigma \cup \{A\} \Vdash B \quad \& \quad \Pi \cup \{C\} \Vdash B \rightarrow \Pi \cup \Sigma \cup \{A \vee C\} \Vdash B$$

$$(R2) \quad \Sigma \Vdash A \Rightarrow B \quad \& \quad \Pi \Vdash A \Rightarrow C \rightarrow \Pi \cup \Sigma \Vdash A \Rightarrow (B \wedge C)$$

$$(R3) \quad \Sigma \Vdash A \Rightarrow B \quad \& \quad \Pi \Vdash C \Rightarrow B \rightarrow \Pi \cup \Sigma \Vdash (A \vee C) \Rightarrow B$$

$$(R4) \quad \Sigma \Vdash A \Rightarrow B \rightarrow \Sigma \Vdash (\neg B \Rightarrow \neg A)$$

$$(R5) \quad \Sigma \Vdash A \quad \& \quad \Pi \Vdash A \Rightarrow B \rightarrow \Pi \cup \Sigma \Vdash B$$

$$(R6) \quad \Sigma \Vdash A \Rightarrow B \ \& \ \Pi \Vdash B \Rightarrow C \rightarrow \Pi \cup \Sigma \Vdash A \Rightarrow C$$

$$(R7) \quad \Sigma \Vdash A \ \& \ \Pi \Vdash B \rightarrow \Pi \cup \Sigma \Vdash A \wedge B$$

$$(R8) \quad \Sigma \Vdash A \Leftrightarrow B \rightarrow \Sigma \Vdash C \Leftrightarrow C(A/B)$$

wobei $C(A/B)$ dadurch aus C entsteht, dass alle Vorkommnisse von B in C durch A ersetzt werden.

$$(R9) \quad \Pi \cup \{A\} \Vdash B \ \& \ \Sigma \Vdash A \rightarrow \Pi \cup \Sigma \Vdash B$$

$$(R10) \quad \Sigma \Vdash P(\acute{a}) \rightarrow \Sigma \Vdash (\exists x)P(x)$$

Existentielle Generalisierung.⁹⁰

$$(R11) \quad \Pi \Vdash (\exists x)P(x) \ \& \ \Sigma \cup \{P(\acute{a})\} \Vdash A \rightarrow \Pi \cup \Sigma \Vdash A$$

wobei $\acute{a}$ weder in $A, (\exists x)P(x)$ noch in einer Annahme in $\Pi \cup \Sigma$ vorkommt.

$$(R12) \quad \Pi \Vdash P(\acute{a}) \rightarrow \Pi \Vdash (\forall x)P(x)$$

wobei $\acute{a}$ in keiner Annahme in Π vorkommt.

$$(R13) \quad \Pi \Vdash (\forall x)P(x) \rightarrow \Pi \Vdash P(\acute{a})$$

$$(R14) \quad \Pi \Vdash \acute{a} = \acute{e} \ \& \ \Sigma \Vdash P(\acute{a}) \rightarrow \Pi \cup \Sigma \Vdash P(\acute{e})$$

Identisches kann für einander eingesetzt werden.

Aus diesen Regeln und Axiomen kann man schon entnehmen:

- (i) **SKP** erfüllt die Extensionalitätsbedingung für die Junktoren " $\neg$ " (so gilt die Doppelte Negation nach (A4) und (A5), der Satz vom Ausgeschlossenen Dritten (A7) und Kontraposition nach (R4)), die Disjunktion (A3) sowie die Konjunktion (nach (A2) und Konjunktionseinführung nach (R7)).
- (ii) **SKP** verwendet ein Enthaltensein, das - wie in der Semantik zu sehen sein muss - stärker ist als bloße Wahrheitsvererbung. Dieses Enthaltensein erfüllt die Modus Ponens-Bedingung (R5).

⁹⁰ vgl. zu diesen Regeln die Regeln von Tennant und die diesbezüglichen Anmerkungen.

(iii) Im Unterschied zu **LP**-Systemen erfüllt **SKP** die Bedingung der minimalen Beschädigung der Konsequenz-Relation besser. So gelten das aussagenlogische Adjunktionsgesetz des Zusammenfassens von Konsequenzen (R2), die Transitivität (R6) und eine Reihe weiterer Regeln wie die Negationseinführung (s.u.). Diesen Regeln entsprechen allerdings keine entsprechenden Theoreme, da Theoreme ja das stärkere Enthaltensein betreffen. Bei **LP**-Systemen verhielt es sich genau anders herum.

Entscheidend ist damit die Frage, ob **SKP** zusätzlich auch die beiden Anti-Trivialitätsbedingungen erfüllt.

Zunächst aber zu einigen **SKP**-Theoremen.

In einer Ableitung werden die Zeilen links numeriert. Rechts neben einer Zeile wird die Berechtigung angegeben, diese Zeile anzuhängen, wobei man sich mittels ihrer Nummern auf vorausgegangene Zeilen bezieht.

Zur besseren Lesbarkeit führen wir folgende Abkürzungen ein:

α, γ statt $\alpha \cup \gamma$ bzw. $\{\alpha\} \cup \{\gamma\}$; zu Beginn einer Ableitung numerieren wir die Annahmen durch und beziehen uns in den (Prämissenmengen der) Sequenzen auf sie mit den betreffenden Nummern - z.B.

(1) $(\forall x)F(x)$	
1. $\Vdash (\forall x)F(x)$	A1
2. $\Vdash F(a)$	R13,1 ⁹¹ ■

⁹¹ Gemäß der Vorschriften für eine Ableitung beziehen sich diese Nummern natürlich immer noch auf die vorangehenden Zeilen, nicht auf die Nummern der Annahmen! Da " $\Vdash$ " nicht das Ableitungszeichen ist, wäre die vollständige Kennzeichnung des Theoremstatus: $\vdash \neg(\forall x(F(x) \wedge G(x))) \Vdash \forall x F(x) \wedge \forall x G(x)$.

(SKPT1) $(\forall x)(F(x) \Rightarrow G(x)), (\forall x)F(x) \Vdash (\forall x)G(x)$

Beweis: (1) $(\forall x)(F(x) \Rightarrow G(x))$

(2) $(\forall x)F(x)$

- | | |
|--|---------|
| 1. 1 $\Vdash (\forall x)(F(x) \Rightarrow G(x))$ | A1 |
| 2. 2 $\Vdash (\forall x)F(x)$ | A1 |
| 3. 1 $\Vdash F(a) \Rightarrow G(a)$ | R13,1 |
| 4. 2 $\Vdash F(a)$ | R13,2 |
| 5. 1,2 $\Vdash G(a)$ | R5,3,4 |
| 6. 1,2 $\Vdash (\forall x)G(x)$ | R12,6 ■ |

Die diesem Theorem entsprechende semantische Folgerungsbeziehung wurde im letzten Abschnitt als für **LPQB** ungültig erwiesen!

(SKPT2) $(\forall x)(F(x) \wedge G(x)) \Vdash (\forall x)F(x) \wedge (\forall x)G(x)$

Beweis: (1) $(\forall x)(F(x) \wedge G(x))$

- | | |
|--|----------|
| 1. 1 $\Vdash (\forall x)(F(x) \wedge G(x))$ | A1 |
| 2. 1 $\Vdash F(a) \wedge G(a)$ | R13,1 |
| 3. $\Vdash F(a) \wedge G(a) \Rightarrow F(a)$ | A2 |
| 4. 1 $\Vdash F(a)$ | R5,2,3 |
| 5. 1 $\Vdash (\forall x)F(x)$ | R12,4 |
| 6. $\Vdash F(a) \wedge G(a) \Rightarrow G(a)$ | A2 |
| 7. 1 $\Vdash G(a)$ | R5,2,6 |
| 8. 1 $\Vdash (\forall x)G(x)$ | R12,7 |
| 9. 1 $\Vdash (\forall x)F(x) \wedge (\forall x)G(x)$ | R7,5,7 ■ |

(SKPT3) $\Vdash A \Leftrightarrow \neg \neg A$

Beweis: aus (A4),(A5),(D1),(R7) ■

(SKPT4) $\Vdash A \wedge A \Leftrightarrow A$

Beweis:

1. $\Vdash A \wedge A \Rightarrow A$ A2
2. $\Vdash A \Rightarrow \neg\neg A$ A4
3. $\Vdash A \Leftrightarrow \neg\neg A$ Theorem 3
4. $\Vdash A \Rightarrow A$ R8,2,3⁹²
5. $\Vdash A \Rightarrow A$ R8,2,3
6. $\Vdash A \Rightarrow A \wedge A$ R2,4,5
7. $\Vdash A \Leftrightarrow A \wedge A$ D1,1,6 ■

d.h. es gilt in **SKP** die Idempotenz der Konjunktion. Zur Erfüllung der starken Anti-Trivialitätsbedingung muss also, nach den Bedingungen für Curry Paradoxien aus Kapitel 3, neben Absorption auch das Modus Ponens-Theorem ungültig sein.⁹³

(SKPT5) $\Vdash (A \wedge B) \vee (A \wedge C) \Rightarrow A \wedge (B \vee C)$

Beweis:

1. $\Vdash A \wedge B \Rightarrow B$ A2
2. $\Vdash B \Rightarrow B \vee C$ A3
3. $\Vdash A \wedge B \Rightarrow B \vee C$ R6,1,2 (Transitivität)
4. $\Vdash A \wedge B \Rightarrow A$ A2
5. $\Vdash A \wedge B \Rightarrow A \wedge (B \vee C)$ R2,3,4
6. $\Vdash A \wedge C \Rightarrow A \wedge (B \vee C)$ analog zu 1-5
7. $\Vdash (A \wedge B) \vee (A \wedge C) \Rightarrow A \wedge (B \vee C)$ R3,5,6 ■

⁹² In (3) wird die Äquivalenz zwischen A und seiner doppelten Negation festgestellt. Also darf nach der Regel der Substitution von Äquivalenten (R8) die doppelte Negation in der linken Seite von (2) durch A substituiert werden.

(SKPT6) $\Vdash (A \wedge B) \vee (A \wedge C) \Leftrightarrow A \wedge (B \vee C)$

Beweis: aus (SKPT5), A6, D1, R7 ■

Es gilt somit die Distributivität von " $\wedge$ " über " $\vee$ ".

(SKPT7) $\Vdash A \Rightarrow A$

Beweis: aus A4, Theorem 3, R8 ■

(SKPT8) $\Vdash A \vee A \Rightarrow A$

Die Idempotenz gilt auch für die Disjunktion.

Beweis: A3, Theorem 7, R3, D1 ■

(SKPT9) $\Vdash (\neg(\neg A \wedge A))$

In **SKP** gilt der Satz vom zuvermeidenden Widerspruch.⁹⁴

Beweis:

1. $\Vdash A \vee \neg A$ A7
2. $\Vdash (\neg(\neg A \wedge \neg \neg A))$ D2,1
3. $\Vdash A \Leftrightarrow \neg \neg A$ Theorem 3
4. $\Vdash (\neg(\neg A \wedge A))$ R8,2,3 ■

Damit läßt sich nun auch als abgeleitete Regel die Widerspruchsregel bzw. die Negationseinführung ableiten, mit der sich *indirekte Beweise* in **SKP** führen lassen:

(AR1) $\Sigma \Vdash A \Rightarrow (\neg B \wedge B) \rightarrow \Sigma \Vdash (\neg A)$

Beweis:

1. $\Sigma \Vdash A \Rightarrow (\neg B \wedge B)$ Voraussetzung (der Regel)
2. $\Sigma \Vdash (\neg(\neg B \wedge B) \Rightarrow \neg A)$ R4,1
3. $\Vdash (\neg(\neg B \wedge B))$ Theorem 9, B/A
4. $\Sigma \Vdash (\neg A)$ R5,2,3 ■

⁹³ Die bei den Bedingungen für Curry Paradoxien genannten Abschlüsse der Theoreme unter Modus Ponens und Substitution von Äquivalenten sind in **SKP** sogar Regeln (nämlich (R5) und (R8)).

⁹⁴ In **LP** gilt dieser natürlich auch, da alle Standard-Tautologien gelten.

Außerdem gelten z.B.:

$$(SKPT10) A \Rightarrow (B \Rightarrow C), A \wedge B \Vdash C$$

ein Analogon zur Importation.

Beweis: A1, A2, R5 ■

$$(AR2) \Sigma \Vdash A \Rightarrow B \ \& \ \Pi \Vdash (\neg B) \rightarrow \Sigma \cup \Pi \Vdash (\neg A)$$

der Modus Tollens als abgeleitete Regel.

Beweis: R4, R5 ■

Wie sieht es nun mit der **Semantik** für **SKP** aus? Für die extensionalen Junktoren können wir - wie gesagt - die Wahrheitstafeln von **LP** übernehmen. Die Semantik von " $\Vdash$ " wurde oben schon erklärt. Ebenfalls übernehmen können wir die Quantorensemantik von **LPQ**. Zu ergänzen sind hier Regelungen für die Identität. Gehen wir wieder von Modellstrukturen aus, wie sie für **LPQ** angegeben wurden, dann legen wir für eine Interpretationsfunktion fest:

$$(SKPS=) (\forall v)(v("=")=id(D))$$

Jede Interpretationsfunktion v irgendeines Modells interpretiert das Gleichheitszeichen als die Identitätsfunktion auf dem Gegenstandsbereich D , d.h. $\{ \langle d, d \rangle \mid d \in D \}$. Die Wahrheitsbedingung für Elementaraussagen findet dann entsprechend auf Identitätsaussagen Anwendung.⁹⁵

Die entscheidende zu regelnde Neuerung, die mit **SKP** auftritt, ist das Enthaltensein, das einen inhaltlichen Relevanten Zusammenhang zwischen Antezedens und Konsequenz ausdrücken soll. Es muss daher eine Semantik für " $\Rightarrow$ " angegeben werden. Daraus ergibt sich mit D1 die Semantik von " $\Leftrightarrow$ ".

⁹⁵ Bei Priest (IC:98) gibt es auch einen Negativbereich für " $=$ ". Dann ergibt sich eine Nicht-Standard-Quantorenlogik, da (A8) nicht mehr gilt. (Darunter kann ich mir nichts vorstellen.)

Priest hat dazu zwei Versuche unternommen:

SEMANTIK (I) Eine Algebra der Bedeutungen⁹⁶

Die erste Option besteht darin, das Enthaltensein als ein Verhältnis zwischen den Bedeutungen von Antezedens und Konsequenz aufzufassen. Das Enthaltensein von B in A wäre dann gegeben, wenn die Bedeutung von B in der Bedeutung von A enthalten ist. Damit läge ein Relevanter Zusammenhang zwischen A und B vor. Die Paradoxien der materialen Implikation können so vermieden werden. Schwieriger ist die Ausgestaltung einer entsprechenden Semantik. Zunächst wäre ein Bereich der Bedeutungen (bezeichnen wir sie hier einfach durch "a", "b", "c" usw.) anzunehmen, der dann durch Postulate bezüglich seiner Struktur gekennzeichnet wird. So fordert Priest bezüglich der Relation des Enthaltenseins der Bedeutung, die durch "<" benannt werden kann:

(IE1) $a < a$

Die Bedeutung von a enthält sich selbst: Reflexivität.

(IE2) $a < b$ und $b < c$, dann $a < c$

Wenn a b enthält und dieses c, so enthält a c: Transitivität.

(IE3) $a < b$ und $b < a$, dann $a = b$

Wenn a b enthält und umgekehrt, dann sind sie identisch: Anti-Symmetrie. Mit diesen Bedingungen ergibt sich im Reich der Bedeutungen eine Halbordnung⁹⁷. Konnexität (die Vergleichbarkeit einer beliebigen Bedeutung mit einer anderen, so dass immer eine in der anderen enthalten sein muss) liegt nicht vor. Deshalb liegt keine Totalordnung vor.

Diese Bedingungen entsprechen noch ohne große Überlegungen unserem intuitiven Verständnis des Enthaltenseins von Bedeutungen.

⁹⁶ vgl. Priest, "Sense, Entailment and *Modus Ponens*".

⁹⁷ zu den Begrifflichkeiten der Algebra und Ordnungstheorie vgl. Stoll, *Set Theory and Logic*, Kap.1 und 6.

(IE3) entspricht (D1), d.h. dem Begriff der logischen Äquivalenz. (IE2) entspricht (R6), und (IE1) entspricht (SKPT7). Für die logischen Beziehungen, die für die Enthaltenseins-Relation gelten sollen - nämlich *alle* Theoreme von **SKP** - müssen jedoch viel mehr Strukturelemente vorhanden sein. Außerdem müssen die Strukturen so sein, dass die ungewünschten Aussagen (wie das Modus Ponens-Theorem) in dieser Struktur keine Entsprechung finden. Dazu entwickelt Priest eine Algebra der Bedeutungen. Grundlegend ist das Prinzip, dass eine Aussage um so mehr bedeutet, je mehr sie über die Welt sagt, d.h. je mehr Sachverhalte sie ausschließt. So kann die Bedeutung von $A \wedge B$ auf keinen Fall mehr zulassen als eine der beiden Teilbedeutungen und

alles, das mindestens so viel Inhalt hat wie A und wie B, hat mindestens so viel Inhalt wie $[A \wedge B]$. In anderen Worten, der Gehalt von $[A \wedge B]$ ist die größte untere Schranke der Inhalte von A und B unter der Enthaltenseins-Ordnung.⁹⁸

Damit werden Strukturen der Algebra für den Bereich der Bedeutungen vorgeschlagen. Und hier drohen uns unsere Intuitionen bezüglich des Enthaltenseins von Bedeutungen ineinander zu verlassen. Ein Verband der Bedeutungen soll sich durch folgende weitere Strukturen auszeichnen⁹⁹:

(IE4) $a \cap b < a$ $a \cap b < b$

Die Bedeutung von " $A \wedge B$ " enthält die Bedeutung jedes der Konjunkte. Das entspricht (A2).

(IE5) $a < b$ und $a < c$, dann $a < b \cap c$

Wenn A die Bedeutung von B und die von C enthält, dann enthält A auch die Bedeutung ihrer Konjunktion. Tut A das? Diese semantische Bestimmung korrespondiert (R2).

(IE6) $a < a \cup b$ $b < a \cup b$

⁹⁸ Priest, "Sense, Entailment and *Modus Ponens*", S.417.

⁹⁹ ebd. S.423f.

Die Bedeutung von $A \vee B$ soll die kleinste obere Schranke bezüglich der Bedeutungen von A und B in der Halbordnung der Bedeutungen sein. Das entspricht (A3).

Entsprechend der doppelten Negation soll gelten, dass der konverse Gehalt a^* von a sich wie folgt verhält:

$$(IE7) \quad a^{**} = a$$

$$(IE8) \quad a < b, \text{ dann } b^* < a^*$$

Dies entspricht der Kontraposition (R4) und macht die Algebra der Bedeutungen zu einem DeMorgan-Verband, jedoch handelt es sich nicht um eine Boolesche Algebra, da aufgrund der antinomischen Aussagen nicht gilt

$$(*IE) \quad a \cap a^* = 0$$

wobei "0" der leere Gehalt wäre. Außerdem soll gelten:

$$(IE9) \quad a \cap (b \cup c) < (a \cap b) \cup (a \cap c)$$

Was soll uns das bezüglich Bedeutungen sagen? Die Bedeutung von $A \wedge (B \vee C)$ soll die Bedeutung von $A \wedge B \vee A \wedge C$ enthalten. Tut sie das? Die Neigung, diese Frage zu bejahen, kann m.E. nur daher rühren, dass wir geneigt sind, auf die Bedeutung der beteiligten logischen Ausdrücke abzuheben, also (A6) für wahr halten. Und zur Bewahrheitung von (A6) wird (IE9) auch postuliert.

Die Wahrheitsbedingung, die sich nun für das Enthaltensein im Rahmen einer parakonsistenten Bewertungsfunktion ergibt, lautet¹⁰⁰:

$$(SI \Rightarrow) \quad 1 \in v(A \Rightarrow B) \leftrightarrow a < b$$

Allerdings verfährt Priest selbst hier mit Filtern auf der Algebra der Bedeutungen, wobei ein Filter sozusagen bestimmte Bedeutungen "herausfiltert". Dieses Herausfiltern soll das Analogon zu einem bestimmten Bewerten von Aussagen sein.

¹⁰⁰ vgl. ebd. S.424f. Zu beachten ist das Priest " $\rightarrow$ " statt " $\Rightarrow$ " verwendet und letzteres Zeichen für eine Enthaltenseinsfunktion eines Modells bezüglich des Bereichs der Bedeutungen verwendet. Das Ergebnis ist jedoch dasselbe.

Um die Bewertung eines Enthaltenseins zu erhalten, müssen allen elementaren Aussagen Bedeutungen zugeordnet werden. Eine Interpretation ordnet nun jeder elementaren Aussage eine Bedeutung zu. Für die anderen aussagenlogisch komplexen Aussagen gilt dann:

$$v(A \wedge B) = v(A) \cap v(B)$$

$$v(A \vee B) = v(A) \cup v(B)$$

$$v(\neg A) = v(A)^*$$

Damit lassen sich nun alle Enthaltenseinsaussagen bewerten nach der Regel $(S \Rightarrow)^{101}$. Dazu muss man sich die dem Antezedens zugeordnete Bedeutung daraufhin ansehen, ob sie die Bedeutung des Konsequenz enthält.

Sollen zugleich die anderen aussagenlogischen Junktoren weiterhin extensional verstanden werden, muss der Filter sich bezüglich der Gültigkeit und der Folgerung so verhalten wie die aussagenlogische Bewertungsfunktion. Das kann man garantieren¹⁰². Priest beweist sogar die Adäquatheit des aussagenlogischen Teils von **SKP** relativ zu dieser Semantik¹⁰³.

Die Schwierigkeiten des ganzen Ansatzes (I) treten indessen auf, wenn nun mit dieser Semantik bestimmte Aussagen als ungültig erwiesen werden sollen. Dazu führt Priest nun nämlich nur algebraische Strukturen als GegenModelle an. So wird das Modus Ponens-Theorem - Priests Beispiel - widerlegt, relativ zu einem Verband der Teilmengen der Menge der natürlichen Zahlen und eines von $\{0\}$ generierten Ultrafilters¹⁰⁴. Nun mag man die Berechtigung dafür darin sehen, dass logische Theoreme unabhängig von ihrem Anwendungsbereich wahr sein sollen, doch war

¹⁰¹ Welche Bedeutungen zugeordnet werden und in welchen Relationen des Enthaltenseins sie stehen (etwa ob die "August ist ein Kater." zugeordnete Bedeutung die Bedeutung "August ist ein Säugetier" enthält), sind Fragen des einzelsprachlichen Bedeutungssystems. Hier werden nur die allgemeinsten Strukturen von solchen Bedeutungssystemen benannt.

¹⁰² Dazu dienen die Bedingungen für Filter bei Priest (ebd. S.424).

¹⁰³ vgl. ebd., S.428f.

der Sinn der ganzen Konstruktion, die Rede von Enthaltensein als Beziehung zwischen den Bedeutungen der Teilaussagen zu explizieren. Relativ zu einem Bereich der Bedeutungen von Aussagen, für den wir intuitiv sicher nicht die Teilmengen der Menge der natürlichen Zahlen adäquat halten, sollte Enthaltensein beurteilt werden. Priests GegenModell zum Modus Ponens-Theorem zuzulassen muss einem deshalb irreführend vorkommen. Versucht man sich hingegen GegenModelle am Begriff des Enthaltenseins von Bedeutungen klarzumachen, gerät man in Schwierigkeiten. Nehmen wir unsere beiden Kandidaten und versuchen sie zu widerlegen:

$$(i) A \wedge (A \Rightarrow B) \Rightarrow B$$

Die Bedeutung von A zusammen mit der Bedeutung, dass A B enthält, muss uns nicht die Bedeutung von B geben.

Das klingt einleuchtend: Das Antezedens macht bezüglich B nur eine Quasi-Metaaussage bezüglich seines Enthaltenseins in A, was immer auch die Bedeutung von B nun genauer sein mag¹⁰⁵.

Weitere Verschachtelungen sind noch schwieriger zu formulieren oder einzusehen:

$$(ii) (A \Rightarrow (A \Rightarrow B)) \Rightarrow (A \Rightarrow B)$$

Wenn es in der Bedeutung von A liegt, dass A B enthält, dann muss es nicht im Sinn *dieser* Aussage liegen, dass A B enthält.

Warum eigentlich nicht? Das scheint das Antezedens doch zu sagen. Oder sagt es bloß, dass in der Bedeutung von A enthalten ist, das A B enthält, während A B nicht *wirklich* enthält? Aber wie kann sich A hier "täuschen"?

Nehmen wir an, wir könnten uns so überzeugen, dass Absorption und Modus Ponens-Theorem nicht gültig sind. Das würde die starke Antitri-

¹⁰⁴ vgl. ebd., S.432.

¹⁰⁵ "Quasi-Metaaussage" meint dabei, dass in der Bedeutung des Antezedens das Enthaltensein (also eine Bedeutungsrelation) gemeint wird. Das erinnert an Freges Problem des "ungeraden Sinns", d.h. des Beziehens auf eine Bedeutungsrelation.

vialität mit sich bringen. Denn die schwache Anti-Trivialität ergibt sich schon daraus, dass die Bedeutung von B gar nichts mit der Bedeutung von $A \wedge \neg A$ zu tun haben muss, also *ex contradictione quod libet* als Enthaltensein nicht gültig ist.

Allerdings reichen die bis jetzt angestellten Überlegungen zur Ungültigkeit der Absorption z.B. nicht aus: Um die Ungültigkeit einer Aussage oder einer Folgerung in der parakonsistenten Logik zu zeigen, reicht es nicht aus, ein Modell zu entwickeln, in dem die Aussage falsch ist - wie es allerdings auch parakonsistente Logiker oft tun -, sondern die Aussage muss vielmehr *ausschließlich* falsch sein. Die bloße Falschheit ließe zu, dass die Aussage auch wahr ist, eben eine Dialetheia, und diese besitzen *per definitionem* einen ausgezeichneten Wahrheitswert. Bezüglich eines Modells muss also gezeigt werden, dass eine fragliche Aussage oder Folgerung *nur* falsch ist. Ob das die obigen Überlegungen zeigen weiß ich nicht.

Vielleicht lässt sich aber eine intuitiv nachvollziehbarere Semantik für **SKP** angeben.

SEMANTIK (II) Eine modale Semantik des Enthaltenseins

Während Priest zu Zeiten der Semantik (I) die Modallogik noch für "bankrott" erklärte¹⁰⁶, legt er in *In Contradiction* eine solche Modalsemantik vor (vgl. IC:102-114).

Zu den möglichen Welten müssen dabei Welten gehören, bei denen bezüglich einer Aussage A sowohl A als auch $\neg A$ wahr sind, wie auch immer sie sich ergeben oder zu konstruieren sind. Ein Modell ist ein Tupel $\langle W, R, @, D, v, G \rangle$: W ist die Menge der möglichen Welten, R die Zugänglichkeitsrelation zwischen diesen Welten, @ ist die aktuelle Welt,

¹⁰⁶ vgl. Priest, "Sense, Entailment and *Modus Ponens*", S.419; interessanterweise bezieht Priest "ähnliche Bemerkungen" (ebd.) ausdrücklich auch auf Routleys Semantik der dreistelligen Zugänglichkeitsrelation.

D ist der Individuenbereich, v eine Interpretationsfunktion, G ist die Menge der Variablenbelegungen. In die Interpretationsfunktion sind jetzt bezüglich der Prädikatoren Weltindices aufzunehmen, das heißt, einem Prädikator wird relativ zu einer Welt ein Positiv- und ein Negativbereich zugeordnet. Für die Individuenausdrücke legen wir hier aus Einfachheitsgründen - da die entsprechende Auseinandersetzung in die modale Prädikatenlogik gehört und nicht in die Debatte um die parakonsistente Logik - fest, dass sie starr designieren (sich immer auf denselben Gegenstand beziehen). Die Wahrheitsbedingungen für die extensionalen Junktoren bleiben erhalten und beziehen sich immer auf Wahrheitswertabhängigkeiten in einer Welt. Nur für das Enthaltensein werden die Weltindices und die Zugänglichkeitsrelation relevant. Die Wahrheitsbedingung lautet:

$$\begin{aligned}
 (SII \Rightarrow) \text{ (i) } 1 \in v(w, A \Rightarrow B) &\leftrightarrow (\forall w' \in W)((w'Rw \ \& \ 1 \in v(w', A) \rightarrow \\
 &1 \in v(w', B)) \ \& \ (w'Rw \ \& \ 0 \in v(w', B) \rightarrow 0 \in v(w', A))) \\
 \text{ (ii) } 0 \in v(w, A \Rightarrow B) &\leftrightarrow (\exists w' \in W)(w'Rw \ \& \ 1 \in v(w', A) \\
 &\ \& \ 0 \in v(w', B))
 \end{aligned}$$

d.h. das Enthaltensein ist wahr in einer Welt w genau dann, wenn in allen Welten, die von w aus zugänglich sind, die Wahrheit von A die Wahrheit von B mit sich bringt und die Falschheit von B die Falschheit von A mit sich bringt; das Enthaltensein ist falsch in einer Welt w , wenn von w aus eine Welt zugänglich ist, in der bei Wahrheit des Antezedens das Konsequenz falsch ist.

Das Enthaltensein ist als modaler Junktor immer noch stärker als bloße Wahrheitsvererbung, doch ist bemerkenswert, dass Enthaltensein in Semantik II nichts mehr zu tun hat mit einer Relation von Bedeutungen wie in Semantik I.

Gültigkeit besagt dann Wahrheit in allen Modellen (bei allen Interpretationen v) *in der aktuellen Welt*, d.h.:

$$(SIIF) \ \Sigma \models \alpha \leftrightarrow (\forall m_i \in M)((\forall \gamma \in \Sigma) 1 \in v_i(@, \gamma) \rightarrow 1 \in v_i(@, \alpha))$$

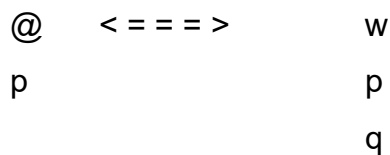
Eine Folgerung ist gültig, wenn in allen Modellen, im Falle, dass die Bewertungsfunktion v alle Prämissen in $@$ wahr macht, sie auch die Konklusion in $@$ wahr macht. Gültig ist dann wieder eine Aussage, die aus einer leeren Prämissenmenge folgt.

Die aktuelle Welt spielt hier die Rolle des Wahrheitsbestimmers. Eine Aussage ist in dieser Semantik logisch wahr, wenn sie in der aktuellen Welt logisch wahr ist. Alle anderen Welten haben nur einen Hilfsstatus. Und dass scheint eine akzeptable Rollenverteilung in einer möglichen Welten-Semantik zu sein.

Das Besondere an dieser Semantik ist die Rolle von $@$ und die Festlegungen bezüglich R . Postuliert wird:

(SIIR) $(\forall w \in W) wR@$

Alle Welten sind von $@$ aus zugänglich. $@$ ist sich selbst zugänglich, aber eine allgemeine Reflexivität der Zugänglichkeit gilt nicht. Wäre $@$ sich nicht selbst zugänglich, würde der Modus Ponens auch *als Regel* nicht gelten, und das soll er ja gerade. Nehmen wir " $< == >$ ", " $= == >$ ", " $< ==$ " als Symbole für gerichtete Zugänglichkeitsrelationen. Betrachten wir folgendes Modell:



In diesem Modell gibt es außer $@$ eine weitere Welt w . Von $@$ kann w erreicht werden. Keine der beiden Welten ist sich selbst zugänglich. Bezüglich einer Bewertung in einem solchen Modell ist anzugeben, welche atomaren Aussagen in einer Welt gelten. Deshalb stehen unterhalb von w " p " und " q " und unterhalb von $@$ steht " p ". Nun ergibt sich aufgrund der Wahrheitsbedingungen der Junktoren folgende weitere Bewertung:

@	===>	w
p		p
		q
p⇒q		

da in w, d.h. allen Welten, die von @ gesehen werden, wenn "p" wahr ist, auch "q" wahr ist. Da aber in @ "q" nicht wahr ist, ist in diesem Modell @ eine Welt, in der der Modus Ponens als Regel nicht gilt, d.h. es gibt ein Gegenmodell, also gilt der Modus Ponens als Regel nicht. Wäre @ sich selbst zugänglich, könnte "p⇒q" nicht als wahr bewertet werden, da trotz der Wahrheit von "p" "q" nicht in @ wahr ist. Es wird also gefordert @R@. Der intuitive Sinn davon ist, dass uns die aktuelle Welt zugänglich ist. Wie könnte das auch anders sein? Diese Annahme erscheint daher harmlos¹⁰⁷. Dass alle anderen Welten von @ aus zugänglich sind, rechtfertigt sich dadurch, dass es sich bei den zu klärenden Möglichkeiten (und anderen Modalbegriffen wie Enthaltensein) um Modalitäten *für uns* bzw. von der aktuellen Welt aus handeln soll. Die Menge der möglichen Welten, die in irgendwelchen Modellen von **SKP** vorkommen, ist also die Menge der Welten, die für uns möglich sind. Stärkere Modalbegriffe treten in der Semantik II nicht auf. Ob sie philosophisch wünschenswert sind, ist auch keine Frage, die die parakonsistente Logik betrifft, sondern wieder eine Frage der Modallogik, die für die hier verhandelten Fragen nicht zuerst beantwortet werden muss.

Absorption ist in dieser Semantik nicht gültig. Gezeigt werden soll, dass Absorption als Folgerungsregel nicht gilt, denn dann gilt sie erst Recht nicht für das stärkere Enthaltensein:

(SIIT1) Nicht: $A \Rightarrow (A \Rightarrow B) \models A \Rightarrow B$

Beweis: In folgendem Modell sei die Zugänglichkeitsrelation bei @ und w2 reflexiv.

¹⁰⁷ Diese Zugänglichkeit bezieht sich nur auf das *Vorliegen* von Tatsachen, nicht auf ein Wissen. Die Reflexivität von @ darf also nicht *mit unserer* Allwissenheit bezüglich der aktuellen Welt gleichgesetzt werden.

@	= = = > w1	= = = > w2
	A	A
		B
	$A \Rightarrow B$	$A \Rightarrow B$ (da w1 nur w2 und w2 nur sich selbst sieht, in @ gilt $A \Rightarrow B$ nicht)
	$A \Rightarrow (A \Rightarrow B)$	(@ sieht w1 und w2)

Obwohl nun $A \Rightarrow (A \Rightarrow B)$ in @ wahr ist, ist $A \Rightarrow B$ nicht wahr, da @ w1 sieht, wo trotz Wahrheit von A B falsch ist. Absorption als Regel gilt also in @ nicht, und ist daher aufgrund der Definition der Folgerung in Semantik II nicht gültig. ■¹⁰⁸

Das Modus Ponens-Theorem ist ebenfalls in dieser Semantik nicht gültig.

(SIIT2) Nicht: $\models (p \wedge (p \Rightarrow q) \Rightarrow q)$

Beweis: Betrachten wir das folgende Modell:

@	< = = = >	w
p		p
q		

wobei @ außerdem sich selbst zugänglich ist. Es ergeben sich folgende weitere Bewertungen:

$p \Rightarrow q$	(da w nur @ sieht)
$p \wedge (p \Rightarrow q)$	(wahrheitsfunktional)

In @ ist " $p \Rightarrow q$ " nicht wahr, da @ w sieht. Nun sieht @ eine Welt, in der " $p \wedge (p \Rightarrow q)$ " wahr ist, aber "q" nicht. Also ist in @ " $p \wedge (p \Rightarrow q) \Rightarrow q$ " nicht wahr

¹⁰⁸ Die weiteren Bewertungen, die in dem Gegen-Modell III auftreten, sind nicht mehr relevant.

(gemäß der Wahrheitsbedingung $(SII \Rightarrow)^{109}$), also ist es - nach Definition der Gültigkeit - nicht gültig. ■

SKP erfüllt also in der Semantik II auch die starke Anti-Trivialitätsbedingung. Außerdem gilt der Modus Ponens als Regel.

Trotzdem bleiben Bedenken: Wir können uns in dieser Semantik Welten aufschreiben oder "vorstellen", die nicht reflexiv sind und keinen Zugang zu allen anderen Welten besitzen (wie gerade in dem Beweis die Welt w).

- Nun, oben wurde eigentlich nur dafür argumentiert, dass von $@$ aus alle Welten zugänglich sein sollten. Bezüglich der anderen Welten scheint daher keine Festlegung nötig, und so ergibt sich

$$(\forall w \in W)(wR@)$$

als das einzige Postulat bezüglich der Zugänglichkeitsrelation. So verfährt Priest in (IC:107f.). Wenn die so von uns vorgestellten Welten aber Weisen sein sollen, wie *unsere* Welt sein könnte. d.h. *Möglichkeiten für uns*, dann scheint das indessen wiederum zu verlangen, dass solche Welten auch sich selbst zugänglich sind und alle anderen Welten betrachten. Denn ist es von uns für uns denkbar, dass diese Zugänglichkeit nicht besteht? Nehmen wir die Zugänglichkeitsargumente bezüglich der aktuellen Welt und die Rede von den mit möglichen Welten vorgestellten Möglichkeiten für uns also ernst, erreichten wir eine universale Zugänglichkeitsrelation zwischen den möglichen Welten, d.h. eine S5-Semantik¹¹⁰. Diesen Weg wählt Semantik (II) jedoch nicht, da in ihr Absorption und das Modus Ponens-Theorem ungültig sein sollen. Insofern befriedigt die modale Semantik (II) für **SKP** auch nicht völlig.

Wenden wir uns nun bezüglich der beiden Semantiken einer metalogischen Frage zu.

¹⁰⁹ Zu sagen, es sei falsch, reicht in einer parakonsistenten Bewertung - wie oben schon gesagt - nicht aus. Die fragliche Aussage muss ausschließlich falsch sein, d.h. "1" darf nicht in ihrer Bewertung vorkommen. Das ist im Beispiel der Fall.

Das wichtigste metalogische Theorem bezüglich eines formalen Systems betrifft seine **Korrektheit**: das, was wir mit dem System beweisen können, muss logisch wahr sein (wenn es aus einer leeren Prämissenmenge folgt) oder muss seine Wahrheit aus der Wahrheit der zugrundegelegten Prämissen erben (Korrektheit der Ableitungsbeziehung). Wünschenswert ist natürlich auch, dass das System deduktiv vollständig ist, d.h. alle logischen Wahrheiten und Folgerungen, die von uns zugrundegelegten Semantik, auch herleiten kann. Entscheidender ist aber die Korrektheit. Wie schon in Kapitel 1 begründet, hängt der Dialethismus ohne die Korrektheit eines Systems der parakonsistenten Logik in der Luft, denn die begründeten Widersprüche sind nur dann wahr, wenn die zur Begründung verwendete Logik korrekt ist. Anders als im Falle der Standard-Logik folgt aus der Korrektheit hier nicht die (einfache) Konsistenz. Parakonsistente Logiken sind ja per definitionem nicht konsistent (d.h. in der metalogischen Sprachregelung: nicht "einfach konsistent"). Was für die Korrektheit zu zeigen ist, ist, dass wir in der betrachteten Logik niemals von wahren Prämissen zu nur falschen Konklusionen geführt werden. Dazu muss gezeigt werden, dass die Axiome logisch wahr sind und die Regeln allein Wahrheit vererben. Für die naive Semantik heißt dies z.B.: Gehen wir von Konvention (T) aus, die wir gemäß der intuitiven Auffassung von Wahrheit für einen Bestandteil der Definition von "wahr" halten, die also logisch wahr ist, dann erhalten wir in einer hinreichend ausdrucksstarken Sprache, die Namen für ihre Aussagen enthält, mit einigen logischen Regeln die Antinomie des Lügners (als Beispiel). Das soll indessen nicht dazu führen, dass irgendwelche falsche Aussagen oder sogar alle Aussagen ableitbar werden (d.h. in der eingeführten metalogischen Redeweise, dass die Theorie "absolut inkonsistent" ist). Auf die Konvention (T) kommt Kapitel

¹¹⁰ Was wieder zeigt, dass die S5-Semantik in unseren Intuitionen bezüglich der Modalbegriffe am besten verankert ist.

7 zu sprechen. Auf die Problematik eines Beweises der Nicht-Trivialität komme ich weiter unten zu sprechen.

LPQ ist ein korrekter Kalkül, so dass der Dialethismus nicht "in der Luft hängt", doch wollten wir ja einen anderen Kalkül, der auch die Modus Ponens-Bedingung voll erfüllt.

Betrachten wir dazu wieder die beiden Semantiken für **SKP**.

(I)

Für die Semantik (I) hat Priest - allerdings unter Verwendung solcher Filtrationsstrukturen, wie oben erwähnt - bewiesen (s.o.), dass der aussagenlogische Teil von **SKP** adäquat, also auch korrekt ist. Zu bedenken sind also (A8), die Identitätsregel (R14) sowie die Quantorenregeln. Die Korrektheit von (A8) ergibt sich trivial aus (SKPS=). Entsprechend gilt (R14), da sich die Wahrheit von $P(\acute{a})$ auf der letztlich zu erreichenden Basisstufe der Elementaraussagen gemäß deren Wahrheitsregel auf den Referenten von $\acute{a}$ bezieht, der, da $\acute{a}=\acute{e}$ wahr ist, identisch sein muss mit dem von $\acute{e}$, so dass auch $P(\acute{e})$ wahr sein muss. Nicht korrekt muss hingegen

(*) $\vdash (\forall x)(\forall y)(x=y \Rightarrow (P(x) \Leftrightarrow P(y)))$

sein. Denn wenn wir die Semantik der Bedeutungsalgebra auf die Bedeutungen von singulären und generellen Termen ausdehnen, dann muss die Bedeutung von $P(\acute{a})$ nicht die Bedeutung von $P(\acute{e})$ enthalten, wenn die singulären Terme eine unterschiedliche Bedeutung haben, obwohl sie sich auf denselben Gegenstand beziehen.

Die Quantorenregeln (R10) - (R12) entsprechen den Quantorenregeln, wie sie sich in Systemen des Natürlichen Schließens finden. Da hier eine Standard-Semantik der Quantifikation zugrundegelegt wird, lässt sich ein entsprechender Korrektheitsbeweis¹¹¹ übertragen. Wir erhalten damit:

¹¹¹ vgl. z.B. Essler/Martinez, *Grundzüge der Logik*, S.325-34.

(SPKMT1) $\Sigma \vdash \alpha \rightarrow \Sigma \models_{\text{SPKS I}} \alpha$

d.h. die Theoreme von **SPK** sind bezüglich der Semantik I korrekt.

(II)

Die quantorenlogischen Überlegungen übertragen sich in die Semantik II, wobei der jetzt konstant gehaltene Gegenstandsbereich aller möglichen Welten die Menge der möglichen Gegenstände sein kann.¹¹² Für den aussagenlogischen Teil von **SPK** ist aber die Korrektheit neu zu zeigen:

(A1) gilt trivialerweise. (A2) - (A7) ergeben sich daraus, dass in den einzelnen Welten die wahrheitsfunktionalen Zusammenhänge der extensionalen Junktoren gelten (also beispielsweise (A2) im Falle der Wahrheit von $A \wedge B$ auch B wahr ist).¹¹³ Die Regeln beziehen sich in ihren Übergängen nur auf die Wahrheitsvererbung, bleiben also "in einer Welt". (R1), (R9) und (R7) sind wahrheitsfunktional gültig. (R4) ergibt sich unmittelbar aus der Wahrheitsbedingung des Enthaltenseins und der Negation. (R5) ist der Modus Ponens, dessen Gültigkeit wir oben diskutiert haben. Die Transitivität (R6) muss gelten, da, wenn in allen Welten, in denen A gilt auch B gilt und in allen Welten, wo B gilt auch C gilt, auch in den A-Welten C gelten muss. Analog argumentierend müssen (R2) und (R3) gelten. Substitution von logisch Äquivalenten (R8) ist aufgrund der Definition logischer Äquivalenz trivialerweise gültig.

Wir erhalten so:

(SPKMT2) $\Sigma \vdash \alpha \rightarrow \Sigma \models_{\text{SPKS II}} \alpha$

¹¹² Diese Rede von Possibilia kann man vermeiden, darauf kommt es hier aber nicht an.

¹¹³ Bei diesem Korrektheitsbeweis muss im Übrigen nur gezeigt werden, dass bei Wahrheit der Prämissen (bzw. des Antezedens) die Konklusion (bzw. das Konsequenz) *mindestens* wahr ist. Selbst wenn es auch noch falsch wäre, wäre dies hier egal, da auch "0,1" ein ausgezeichnete Wahrheitswert ist! Im Gegensatz zu den üblichen Korrektheitsbeweisen muss also nicht die Falschheit der Konklusion ausgeschlossen werden. Allerdings muss dafür der Korrektheitsbeweis um einen Anti-Trivialitätsbeweis ergänzt werden. In diesem ist zu zeigen, dass relativ zur zugrundgelegten Semantik mindestens eine Aussage ausschließlich falsch ist.

Die Theoreme von **SKP** sind bezüglich der Semantik II korrekt. Nun haben wir oben mit (SIIT2) bewiesen, dass das Modus Ponens-Theorem nicht *gültig* in Semantik II ist. Wäre es trotzdem in **SKP** ableitbar, wäre **SKP** nicht korrekt. **SKP** ist aber - wie gerade gesehen - korrekt. Also ist das Modus Ponens-Theorem nicht ableitbar. Also gibt es (mindestens) eine Aussage, die sich in **SKP** nicht herleiten lässt.■:

(SKPMT3) $(\exists \alpha)(\text{Nicht: } \vdash_{\text{SKP}} \alpha)$

d.h. **SKP** ist absolut konsistent. **SKP** ist also bewiesenermassen nicht trivial.

Es fehlt jedoch ein Vollständigkeitsbeweis für **SKP**.¹¹⁴

LP besitzt in **LPB** ein Entscheidungsverfahren. **LPQB** ist ein vollständiger Kalkül, wenn er wohl auch als Prädikatenlogik kein effektives Entscheidungsverfahren besitzt¹¹⁵. Bezüglich der Semantik II von **SKP** lässt sich ein Analogon zu einem modallogischen Entscheidungsverfahren entwickeln. In modallogischen Entscheidungsverfahren (für Systeme wie **T**, **S4**, **S5** usw.) führen die Bewertungen der Modalausdrücke zu "neuen" sichtbaren Welten¹¹⁶. Der einzige Modalausdruck in **SKP** ist " $\Rightarrow$ ". Wir brauchen also Regeln zur "Welterzeugung" relativ zu den Bewertungsmöglichkeiten von " $\Rightarrow$ ". Als Verfahren kann festgelegt werden:

1. Bewerte die zu falsifizierende Aussage in @ mit "T'".
2. Verwende in einer Welt w die LPQB-Regeln mit Ausnahme von " $\supset$ "-Regeln.

¹¹⁴ Die üblichen Verfahren zum Beweis der Vollständigkeit stehen in der parakonsistenten Logik nicht zur Verfügung, da diese über den Begriff der maximal *konsistenten* Mengen ansetzen (Henkin-Beweise) oder vom Deduktionstheorem und indirekten Beweisen Gebrauch machen.

¹¹⁵ Diese abgeschwächte Formulierung ergibt sich aus den Konsequenzen, die der parakonsistente Ansatz für die Theoreme der Metalogik hat. Die Beweise, die in der Standard-Logik für die Unentscheidbarkeit der Prädikatenlogik gegeben werden (etwa über das "Halteproblem") lassen sich nicht direkt auf parakonsistente Logiken übertragen!

¹¹⁶ vgl. Hughes/Cresswell, *Introduction to Modal Logic*, S.82-122.

3. Angenommen " $\Rightarrow$ " ist der Hauptjunktork einer Aussage, die in einer Welt w bewertet wird¹¹⁷:

(i) Ist in w $T(A \Rightarrow B)$ so muss in allen Welten, die w sehen kann, TB der Fall sein, wenn TA oder $F'A$ der Fall ist.

(ii) Ist in w $T'(A \Rightarrow B)$, so füge zu den zugänglichen Welten eine hinzu und in dieser Welt ist TA und $T'B$ der Fall.¹¹⁸

(iii) Ist in w $F(A \Rightarrow B)$, so füge zu den zugänglichen Welten eine hinzu und in dieser Welt ist TA und FB der Fall.

(iv) Ist in w $F'(A \Rightarrow B)$, so muss in allen zugänglichen Welten, in denen TA der Fall ist, $F'B$ der Fall sein.¹¹⁹

Angenommen in einer Aussage kommt " $\parallel$ " mit Prämissen A_i vor, dann gilt innerhalb einer betreffenden Welt:

(i) Bei $T(A_i \parallel B)$ oder $F'(A_i \parallel B)$ ist $(T'A_i$ oder $TB)$ der Fall.¹²⁰

(ii) Bei $T'(A_i \parallel B)$ oder $F(A_i \parallel B)$ ist Ta_i und $T'B$ der Fall.¹²¹

4. Alle Welten müssen von $@$ aus zugänglich sein.

5. Tritt in einer Welt eine inkompatible Bewertung (vgl. Kapitel 4.1) auf, so ist die zu falsifizierende Aussage SKP-gültig.

Betrachten wir als Beispiel die nun (im Gegensatz zum kompletten System **SKP**) mögliche Widerlegung der Kontraposition (im restringierten System **SKP**⁻, ohne (R4))¹²²:

¹¹⁷ Die folgenden Bedingungen korrespondieren ($SII \Rightarrow$). Bei diesen Regeln wird die zweite Klausel in ($SII \Rightarrow$)(i) fallengelassen, die der Kontraposition entspricht.

¹¹⁸ Dies ergibt sich aus der Negation der Allaussage ($SII \Rightarrow$)(i).

¹¹⁹ Dies ergibt sich aus der Negation der Existenzaussage ($SII \Rightarrow$)(ii)

¹²⁰ Dies ergibt sich aus ($S \parallel$)(i). Da eine Aussage mit " $\parallel$ " nicht antinomisch sein kann (s.o.), treten eigentlich nur zwei Fälle auf.

¹²¹ Das ergibt sich aus ($S \parallel$)(ii).

¹²² Entsprechende GegenModelle gelten für verschiedene Formen der Negationseinführung wie $(A \Rightarrow \neg A), \neg(\neg A \Rightarrow A) \parallel (\neg A)$.

@	$T'(A \Rightarrow B \parallel (\neg B \Rightarrow \neg A))$	
	$T(A \Rightarrow B), T'(\neg B \Rightarrow \neg A)$	(Regel " $\parallel$ "(ii))
	↓	(Sehen einer Welt nach Regel " $\Rightarrow$ "(ii))
w2	$T\neg B, T'\neg A$	
	$F'A, FB$	(Regel " $\neg$ ")
	TB	(Regel " $\Rightarrow$ "(i))

Die Behauptung, dass $\neg B \Rightarrow \neg A$ in @ nur falsch ist, führt zur Sichtbarwerdung einer weiteren Welt w2, wobei jedoch keine inkompatiblen Bewertungen auftreten, d.h. dieses Modell ein Gegenmodell zur Kontraposition liefert. Die Regel für

$T(A \Rightarrow B)$ fordert keine weiteren sichtbaren Welten, sondern stellt nur eine Bedingung an die vorliegenden Welten, die aber hier nicht zu einer inkompatiblen Bewertung führt. ■

Da es sich hier um eine quantorenlogische Sprache handelt, liegt kein effektives Entscheidungsverfahren vor. Trotzdem handelt es sich um ein Verfahren, so dass

(SKPMT4) (i) nicht($\models_{SKP} \alpha$) $\rightarrow$ α ist falsifizierbar

(ii) α ist nicht falsifizierbar $\rightarrow \models_{SKP} \alpha$

" α ist falsifizierbar" heißt, im Welten-Tableau treten nicht inkompatible Bewertungen auf. Begründung des Meta-Theorems: Ist α nicht falsifizierbar, treten im Tableau, dessen Regeln den semantischen Regeln korrespondieren, beim Versuch der Falsifikation inkompatible Bewertungen für irgendeine Aussage auf, d.h. es gibt *nicht* eine Bewertung v , die α ausschließlich falsch macht, also muss α einen ausgezeichneten Wahrheitswert in allen Bewertungen haben. Ist α hingegen nicht gültig, muss es nach den Wahrheitsbedingungen der in α vorkommenden Ausdrücke,

die sich in den Tableau-Regeln, die logische Komplexität abbauen, spiegeln, eine Bewertung geben, die α nur falsch macht, und diese falsifiziert α . ■

Wir in den jeweiligen Abschnitten festgestellt, weisen beide SKP-Semantiken Mängel auf. Semantik I ist allgemein schwer nachzuvollziehen. Semantik II liberalisiert die Zugänglichkeitsrelation zwischen möglichen Welten stärker als das philosophisch akzeptabel erscheint. Eine Semantik für **SKP**, die beide Mängel überwindet steht daher noch aus.

Das Problem der Hyper-Widersprüche

Wir haben gesehen, wie gegen mehrwertige Lösungsvorschläge zum Antinomienproblem verstärkte Lügner konstruiert werden können. Das gilt für die parakonsistente Lösung nicht. Der verstärkte Lügner bei mehrwertigen Ansätzen, zu denen man die Parakonsistenz aufgrund ihrer dritten Verhaltensmöglichkeit von Aussagen in gewissem Maße rechnen kann, hatte die Form:

(1) Aussage (1) ist nicht wahr.

wobei "nicht-wahr" alle Werte jenseits des Wahren umfaßt. Nun, in einer parakonsistenten Semantik ist (1) einfach eine Aussage, die gemäß der Auswertung von (1), wie wir sie in Kapitel 1 vorgenommen haben, eine weitere Dialetheia mit sich bringt. Dasselbe gilt für die Aussage:

(2) Aussage (2) ist nur falsch.

Diese Aussage nimmt die Unterscheidung von "mindestens wahr", "mindestens falsch", "allein falsch", "allein wahr" zum Ausgangspunkt. Formalisiert besagt (2):

(3) $\text{Wahr}(2) \Leftrightarrow (\text{Falsch}(2) \wedge \neg \text{Wahr}(2))$

Außerdem gilt semantisch (z.B. in **LP**):

(4) $\neg \text{Wahr}(\alpha) \Rightarrow \text{Falsch}(\alpha)$

aufgrund der Einbeziehung des Standard-Verhaltens der Negation. Nun gilt, dass die Auswertung von (2) bzw. (3) zu

$$(5) \text{Wahr}(2) \wedge \neg \text{Wahr}(2)$$

führt, also zu einer Dialetheia.

Beweis:	1. Wahr(2)	AE
	2. Falsch(2) $\wedge$ $\neg$ Wahr(2)	$\supset$ B,(3),1
	3. $\neg$ Wahr(2)	$\wedge$ B,2
	4. Wahr(2) $\wedge$ $\neg$ Wahr(2)	$\wedge$ E,1,3
	5. $\neg$ Wahr(2)	$\neg$ E,1,4
	6. $\neg$ Wahr(2)	AE
	7. $\neg$ (Falsch(2) $\wedge$ $\neg$ Wahr(2))	Kontraposition, $\supset$ B,6,(3)
	8. $\neg$ Falsch(2) $\vee$ $\neg$ $\neg$ Wahr(2)	DeMorgan,7
	9. $\neg$ Falsch(2) $\vee$ Wahr(2)	$\neg$ B,8
	10. $\neg$ Falsch(2) $\supset$ $\neg$ $\neg$ Wahr(2)	Kontraposition von (4)
	11. $\neg$ $\neg$ Wahr(2) $\vee$ Wahr(2)	Dilemma,9,10
	12. Wahr(2)	$\neg$ B,Idempotenz $\vee$,11
	13. Wahr(2) $\wedge$ $\neg$ Wahr(2)	$\wedge$ E,12,5 ■

Dabei sind alle Schritte auch in **SKP** gültig¹²³.

Diese Formen von verstärkten Lügner zeigen also nur, dass es nicht nur einen sondern mehrere Lügner gibt. Alle diese selbstbezüglichen Lügner führen zu Dialetheias. Das heißt, es gibt eine Reihe von wahren Widersprüchen, dies jedoch ist kein grundsätzlicher Einwand (wie im Falle der verstärkten Lügner gegen mehrwertige Ansätze) sondern nur ein weiterer Effekt der Ausdrucksstärke einer universalen Sprache. Insbesondere lassen sich ausgehend von einem Lügner sofort weitere wahre Widersprüche ableiten:

¹²³ DeMorgan ist mit (D2), Dilemma mit (R1) und (A3) gegeben.

Angenommen (5) sei eine der Lügner-Aussagen. Sie hat die Form $\alpha \wedge \neg \alpha$. Nun gilt aber in **LPQ** oder **SKP** der Satz vom Ausgeschlossenen Widerspruch. Wir erhalten so aus einer Lügneraussage A und dem Theorem $\neg(A \wedge \neg A)$ per Konjunktionseinführung:

$$(6) (A \wedge \neg A) \wedge \neg(A \wedge \neg A)$$

Diese Aussage hat wieder die Form eines Widerspruchs und ist selbst eine Dialetheia. Mit Hilfe ihrer ließe sich nun wieder eine Dialetheia beweisen - usw. Mit einem Widerspruch sind also beliebig viele Widersprüche gegeben, doch haben diese alle dieselbe Form und führen weder zur Ableitbarkeit beliebiger Aussagen noch zu andersartigen Widersprüchen.

Insbesondere treten in einem parakonsistenten Ansatz die weiteren Lügner nicht in einer *Metasprache* auf (wie im Fall der mehrwertigen Ansätze). Wäre die parakonsistente Semantik in einer Metasprache formuliert, wäre von dieser z.B. zu fragen, ob sie konsistent ist. Die Überwindung der Objekt-/Metaspracheunterscheidung ist indessen gerade das Anliegen der starken parakonsistenten Semantik. Insofern also die Objektsprache inkonsistent ist, ist - per Identität - auch die (vermeintliche) Metasprache inkonsistent. Es gibt daher keinen weiteren Argumentationsschritt bezüglich einer getrennten Metasprache.¹²⁴

Priest selbst führt ein anderes Problem an¹²⁵:

Wie verhält es sich mit der Wahrheit der Aussage "A ist wahr", wenn A eine Dialetheia ist? In einer mehrwertigen Logik kann es sein, dass Aussagen, die über Wahrheitswerte von Aussagen sprechen, sich bivalent verhalten: "A ist wahr" ist nur dann wahr, wenn A wahr ist, denn wenn A

¹²⁴ Insofern liegt Brendels Priest-Kritik, die über verstärkte Lügner ansetzt und nach der Metasprache fragt (vgl. Brendel, *Die Wahrheit über den Lügner*, S.201ff.), etwas schief. Außerdem haben wir gerade gesehen, dass die Argumentation zum verstärkten Lügner (3) in der parakonsistenten Logik durchführbar ist (vs. ebd. S.203).

¹²⁵ vgl. Priest, "The Logic of Paradox", S.238f.

unbestimmt ist, ist "A ist wahr" einfach falsch¹²⁶. Was, wenn A eine Dialetheia ist? Priest behauptet nun, dass dann "A ist wahr" selbst eine Dialetheia ist! Wir hätten dann folgende Wahrheitstafel:

A	A ist wahr
1	1
0	0
0,1	0,1

und entsprechend für "A ist falsch".

Wenn dem so wäre, hätte das paradoxe Konsequenzen. Die Hauptthese der parakonsistenten Logik ist, dass einige Widersprüche wahr sind. Nun, nach obiger Tafel hat diese These die Form "A ist wahr" bezüglich einer (oder mehrerer) Dialetheia. Damit ist diese Aussage selbst eine Dialetheia. Die Hauptthese der parakonsistenten Logik wäre also eine Antinomie! Priest - mit seiner Vorliebe für das Paradoxe - übernimmt diese Konsequenz¹²⁷. Das ist m.E. absurd:

Ein Mindestbedingung an eine philosophische These sollte lauten, dass sie beansprucht, ausschließlich wahr zu sein. Ansonsten träte die Trivialität, die bezüglich von Theorien gerade vermieden werden sollte, bei ihr auf. Denn eine Antinomie behauptet nichts. Deswegen waren z.B. die verschiedenen Lügner logisch zu isolieren. Es mag zwar Aussagen geben, die zugleich wahr und falsch sind, es gibt aber keinen Anlass, sie in einer theoretischen Auseinandersetzung zu behaupten, da mit ihnen nichts ausgeschlossen wird. Der Sprechakt des Behauptens, mit dem ein Sprecher gegenüber anderen (i) etwas durch die Angabe oder das Verfügen über Gründe als vorliegend hinstellt und (ii) diese Information als äußerungswürdig betrachtet, da die Anerkennung ihrer Wahrheit einen

¹²⁶ so z.B. Blau, *Die dreiwertige Logik der Sprache*, S.82ff.

¹²⁷ vgl. Priest, "The Logic of Paradox", S.238ff.; diese Passage lautet auch "Abschließende selbstbezügliche Nachschrift".

Unterschied in der weiteren Debatte machen würde, kann bezüglich einer Dialetheia nur fehlschlagen. Das Ziel des Behauptens ist allein die Wahrheit. Wenn wir etwas behaupten, dann sind wir auf Wahrheit aus, genauer: auf ausschließliche Wahrheit. Das machen wir uns gewöhnlich nicht klar, da wir uns meist in nicht-antinomischen Kontexten bewegen. In diesen reicht es aus, die Falschheit der Behauptung des Opponenten zu zeigen, um reine Wahrheit zu erreichen. Für Antinomien hingegen können wir im gewöhnlichen Sinne nicht argumentieren, da hier das Gelingen oder Versagen einer Argumentation irrelevant ist für das Einlösen eines Behauptungsanspruches. Eine von uns gelieferte Begründung spielt direkt dem Opponenten in die Hände. Antinomien können wir nicht behaupten *mit dem Wissen*, daß es sich um eine Antinomie handelt, da ihre Behauptung irrelevant ist gegenüber der Behauptung ihres Gegenteils. Antinomien verstoßen gegen die Konversationsmaxime, nur relevante Aussagen zu behaupten und dadurch ihr Gegenteil auszuschließen. Behaupten geht über wahres Meinen plus irgendeine Begründung Besitzen hinaus. Die Handlung des Behauptens macht keinen Sinn bei Antinomien. Bewiesene Antinomien können wir, die wir den Beweis kennen, also nicht mehr behaupten, da dies unserer Praxis des Behauptens zuwiderläuft. Deshalb darf der Dialethismus als philosophische Position auch nicht selbst antinomisch sein. Selbst wenn einige Widersprüche wahr und beweisbar sind, können wir sie doch nicht behaupten. Auch die Gültigkeit des Satzes vom vermeidenden Widerspruch in parakonsistenten Logiken bringt nicht mit sich, dass die These des Dialetheisten eine Antinomie ist: Aus

$$(1) \models (\neg(A \wedge \neg A))$$

ergibt sich aus der Definition der Gültigkeit die Feststellung:

$$(2) T(\neg(A \wedge \neg A))$$

Daraus ergibt sich gemäß der Wahrheitstafel der Negation:

$$(3) F(A \wedge \neg A)$$

Das Argument des Aristotelikers (der behauptet, dass es keine wahren Widersprüche gibt) müsste nun wie folgt weiterlaufen:

$$(4) \neg T((A \wedge \neg A))$$

$$(5) \forall A(\neg T(A \wedge \neg A))$$

da A beliebig gewählt worden war. Also mit Quantorendualität:

$$(6) \neg(\exists A)T(A \wedge \neg A)$$

was man als die These des Aristotelikers ansehen könnte. Der Schritt von (3) zu (4) indessen ist parakonsistent gerade nicht gültig, da man aus der Falschheit einer Aussage nicht darauf schließen kann, dass sie nicht (auch) wahr ist. Die Negation der These des Dialetheisten (d.h.(6)) lässt sich somit nicht einfach aus dessen Zugestehen des Nichtwiderspruchsprinzips gewinnen.

Meiner Meinung nach sollte der Dialetheist daher sagen: Antinomien sind wahr, aber nicht behauptbar. Bezüglich des Operators oder Prädikators "Behaupten" würden dann nicht gelten:

$$(*) \text{ Behaupten}(T(A)) \supset \text{Behaupten}(A)$$

Selbst wenn man die Wahrheit von A behaupten kann (etwa, weil es einen Beweis für A gibt, wie es Beweise für Antinomien gibt), impliziert dies dennoch nicht, dass man A behaupten kann. Dies liegt m.E. - wie oben angedeutet - daran, dass Antinomien keinen kognitiven Gehalt haben. Aus meta-logischen Gründen mag ich als Dialetheist motiviert sein, Widersprüche als wahr zu behaupten, ich behaupte aber nicht unmittelbar eine Antinomie wie den Lügner. Denn was würde ich damit sagen?

Eine Argumentation in diese angedeutete Richtung wäre eine Lösung des Problems im Bereich der Pragmatik (etwa mittels einer Theorie des Behauptens und entsprechender Prinzipien eines Behauptens-Operators). Da diese Lösung ausgearbeitet nicht vorliegt, empfiehlt es sich, das Entstehen des Problems schon in der Semantik zu beheben,

d.h. in der Semantik zu verhindern, dass die Behauptung der Wahrheit einer Antinomie selbst antinomisch ist.

Woher rührt also das Problem? Es ergibt sich ausgehend von der Konvention (T), die bei Aufhebung der Unterscheidung von Objekt- und Metasprache lautet:

$$(T) A \text{ ist wahr} \Leftrightarrow A.$$

Per Kontraposition ergibt sich daraus:

$$(T') \neg A \Leftrightarrow \neg(A \text{ ist wahr})$$

Ist A nun eine Dialetheia, dann gilt $0 \in v(A)$, also gemäß der Wahrheitstafel für die Negation $1 \in v(\neg A)$. Und so mit Modus Ponens aus (T')

$$1 \in v(\neg(A \text{ ist wahr}))$$

also per Wahrheitstafel der Negation:

$$0 \in v(A \text{ ist wahr}).$$

Außerdem gilt dass A wahr ist, also zusammen

$$0 \in v(A \text{ ist wahr}) \text{ und } 1 \in v(A \text{ ist wahr}).$$

"A ist wahr" ist dann selbst eine Dialetheia.

Wie ließe sich das Argument unterbrechen, wenn wir hier zunächst die Konvention (T) als sakrosant betrachten? Modus Ponens können wir nicht fallenlassen, ohne die Modus Ponens-Bedingung zu verletzen. Die Wahrheitstafel der Negation können wir nicht einfach ändern, ohne die Extensionalitätsbedingung zu verletzen. Der Schuldige ist die Kontraposition.

Kontraposition in der Konvention (T) führt dazu, dass $T(A)$ bezüglich einer Antinomie A sowohl wahr als auch falsch ist. Dies gilt unabhängig davon, ob wir " $T()$ " in (T) als Operator oder als Prädikator auffassen.¹²⁸

¹²⁸ " $v(A,1)$ " hingegen drückt aus, dass eine Aussage mittels ihres *Namens* von der Bewertungsrelation auf einen Wahrheitswert bezogen wird. In " $v(, ')$ " kann man nicht mittels (T) substituieren. Ansonsten ergäbe sich selbst bei Preisgabe der Kontraposition die Schwierigkeit. Dann müßte sogar das Verständnis der Negation modifiziert werden.

Ohne Kontraposition der Konvention (T) ließe sich aus der Falschheit von A, also der Wahrheit von $\neg A$, *nichts* für den Wahrheitswert von "A ist wahr" ableiten. Und das soll in einer parakonsistenten Semantik auch so sein, da die beiden Wahrheitswerte "wahr" und "falsch" die Unabhängigkeit von einander genießen müssen, die parakonsistente Bewertungen erlaubt. Wie sieht es nun aus mit der Kontraposition in unseren beiden Systemen **LP** und **SKP**?

In **LP** gilt die Kontraposition als Folgerung. Beweis:

1. $T(A \supset B)$			
2. $T'(\neg B \supset \neg A)$			
3. $F'(\neg B)$		$T' \supset, 2$	
4. $T'(\neg A)$		$T' \supset, 2$	
5. $T'B$		$F' \neg, 3$	
6. $F'A$		$T' \neg, 4$	
7a. FA	$T \supset, 1$		7b. TB $T \supset, 1$
8a. #	7a, 6		8b. # 7b, 5 ■

Beide Äste des Baums schließen, die überprüfte Folgerung ist also LP-gültig. Das entsprechende Konditional gilt als Standard-Tautologie sowieso in LP.

Es ist auch nicht zu sehen, wie sich die Gültigkeit der Kontraposition in **LP** vermeiden ließe. Damit ist m.E. **LP** endgültig erledigt: nicht nur ist der Modus Ponens als Regel nicht gültig, vielmehr ergibt sich unter Verwendung von **LP** auch, dass die These der Parakonsistenz selbst eine Antinomie ist. Das ist inakzeptabel.¹²⁹

¹²⁹ Priest sagt an einer Stelle, dass er die im "Selbstbezüglichen Nachwort" verwendete Kontraposition bezüglich Konvention (T) nun *ausdrücklich* ablehnt (LT:255), in (IC) tritt dies mehr als eine *Option* auf. In einem wenig später erschienen Aufsatz lehnt er dann eine Kritik mit dem Hinweis auf die Logik **LP** ab, in der Modus Ponens nicht gilt - aber eben die Kontraposition - und die als Logik in (LT) angegeben wird! Nun könnte er zwar meinen, die parakonsistente Logik sei LP und in der Konvention (T) komme ein anderes Konditional vor, doch dann wäre die Logik dieses anderen Konditionals zu klären und hier könnten die Gegenargumente wieder einsetzen. Priest Unklarheit

Wie sieht es in **SKP** aus? Kontraposition ist dort sogar eine Regel, nämlich (R4). Warum war diese Regel korrekt? Aufgrund einer diesbezüglichen ausdrücklichen Festlegung in den Semantiken für **SKP**. So lautete die Wahrheitsbedingung des Enthaltenseins in Semantik II:

$$\begin{aligned}
 (\text{SII} \Rightarrow) \text{ (i) } 1 \in v(w, A \Rightarrow B) &\leftrightarrow (\forall w' \in W)(w'Rw \ \& \ (1 \in v(w', A) \rightarrow \\
 &1 \in v(w', B)) \ \& \ (0 \in v(w', B) \rightarrow 0 \in v(w', A))) \\
 \text{ (ii) } 0 \in v(w, A \Rightarrow B) &\leftrightarrow (\exists w' \in W)(w'Rw \ \& \ 1 \in v(w', A) \\
 &\ \& \ 0 \in v(w', B))
 \end{aligned}$$

Das zweite Konjunkt in Klausel (i) garantiert die Gültigkeit der Kontraposition. Lassen wir es fallen, gibt es Gegenmodelle zur Kontraposition:

@	= =>	w
A		A
B		B
		$\neg B$

Die aktuelle Welt @ sieht sich und eine *inkonsistente Welt* w. Solche Welten sind in der parakonsistenten Logik nicht nur zugelassen, sondern müssen ja auftreten, da die Dialetheias wahr sind¹³⁰. In diesem Modell ist $A \Rightarrow B$ wahr, da alle A-Welten, die von @ aus zugänglich sind, auch B-Welten sind. Es gibt aber eine $\neg B$ -Welt, die nicht zugleich eine $\neg A$ -Welt ist. Also ist $(\neg B \Rightarrow \neg A)$ für @ nicht wahr, sondern falsch. Also ist Kontraposition nach (S $\Rightarrow$) *nur* falsch für @, also in der modifizierten Semantik II ungültig. Die Lösung unseres Problems besteht also in einer modifizierten Semantik II, mit der komplexen¹³¹ Wahrheitsbedingung

darüber, was die parakonsistente Logik genau ist (d.h. welches festgelegtes System) dient eher der Verwirrung - nicht nur der Kritiker.

¹³⁰ Außerdem könnten sie nützen, Meinungssysteme von Akteuren zu Modellieren, da diese durchaus Widersprüchliches glauben mögen.

¹³¹ Anstatt einfach zwei " $\leftrightarrow$ " zu verwenden bezüglich des hier immer noch quasi-metasprachlich verwendeten " $\rightarrow$ ", müßten wir mit der Preisgabe der Unterscheidung von Objekt- und Metasprache auch hier auf die Möglichkeit der Kontraposition verzichten. Deshalb müssen alle Verhältnisse des Enthaltenseins - auch die, die sich ansonsten durch Kontraposition ergäben - eigens aufgeführt werden. Wir erhalten so

- (SII $\Rightarrow$) (i) $1 \in v(w, A \Rightarrow B) \rightarrow (\forall w' \in W)(w'Rw \ \& \ 1 \in v(w', A) \rightarrow 1 \in v(w', B))$
- (ii) $\neg(1 \in v(w, A \Rightarrow B)) \rightarrow \neg(\forall w' \in W)(w'Rw \ \& \ 1 \in v(w', A) \rightarrow 1 \in v(w', B))$
- (iii) $(\forall w' \in W)(w'Rw \ \& \ 1 \in v(w', A) \rightarrow 1 \in v(w', B)) \rightarrow 1 \in v(w, A \Rightarrow B)$
- (iv) $\neg(\forall w' \in W)(w'Rw \ \& \ 1 \in v(w', A) \rightarrow 1 \in v(w', B)) \rightarrow \neg(1 \in v(w, A \Rightarrow B))$
- (v) $0 \in v(w, A \Rightarrow B) \rightarrow (\exists w' \in W)(w'Rw \ \& \ 1 \in v(w', A) \ \& \ 0 \in v(w', B))$
- (vi) $\neg(0 \in v(w, A \Rightarrow B)) \rightarrow \neg(\exists w' \in W)(w'Rw \ \& \ 1 \in v(w', A) \ \& \ 0 \in v(w', B))$
- (vii) $(\exists w' \in W)(w'Rw \ \& \ 1 \in v(w', A) \ \& \ 0 \in v(w', B)) \rightarrow 0 \in v(w, A \Rightarrow B)$
- (viii) $\neg(\exists w' \in W)(w'Rw \ \& \ 1 \in v(w', A) \ \& \ 0 \in v(w', B)) \rightarrow \neg(0 \in v(w, A \Rightarrow B))$

wobei die Regel (R4) fallengelassen werden muss. Damit gehen allerdings auch die Negationseinführung und der Modus Tollens verloren!¹³² Die stark anti-triviale parakonsistente Logik wird also immer schwächer.¹³³

Nicht alle Widerlegungen einer Annahme sind jedoch unmöglich:

für (SII $\Rightarrow$)' statt zwei Klauseln acht Klauseln. Ich führe hier einmal eine solche komplexere Definition vor, werde jedoch aus Abkürzungsgründen gelegentlich wieder " $\leftrightarrow$ " so verwenden, dass sich dahinter alle diese Bedingungen verbergen, wobei dies dann eigens festgestellt wird, um Mißverständnisse zu vermeiden. All diese Bemerkungen gelten natürlich nicht für ein LP-Konditional.

¹³² Die für die Transzendentalpragmatik konstitutive aber nie formal durchgeführte oder plausibilisierte These, dass die Umgangssprache gegenüber der Antinomienproblematik immer auch die letzte Metasprache ist (vgl. z.B. Apel, *Transformation der Philosophie*, I, 156ff., 310-18 und II, 247f., 302f., 323f. u.ö), kann zwar durch die parakonsistente Logik eingeholt werden, doch geht dabei gerade das für die Transzendentalpragmatik konstitutive Verfahren des *indirekten* Beweises (z.B. durch Negationseinführung; vgl. Kuhlmann, *Reflexive Letztbegründung*, S.23, 91ff.) verloren. Die Transzendentalpragmatik als "Letztbegründung" hängt also immer noch in der Luft.

¹³³ Auch läßt sich obiges Argument bezüglich (3) nicht mehr durchführen. Das scheint mir aber eine harmlose Konsequenz zu sein.

(AR2) $A \Vdash (\neg A) \rightarrow \Vdash (\neg A)$

Beweis:	1. $A \Vdash (\neg A)$	AE
	2. $\neg A \Vdash (\neg A)$	Axiom 1
	3. $A \vee \neg A \Vdash (\neg A)$	R1,1,2
	4. $\emptyset \Vdash (A \vee \neg A)$	Axiom 7 (zur Verdeutlichung mit " $\emptyset$ ")
	5. $\emptyset \Vdash (\neg A)$	R9,3,4 ■

Entsprechend gilt auch:

(AR3) $\Vdash A \Rightarrow \neg A \rightarrow \Vdash (\neg A)$

Beweis: mit (SKPT7), R3, Axiom 7, R5 ■

d.h. dass im Falle, dass die Wahrheit von A die Wahrheit von $\neg A$ nach sich zieht oder A sogar $\neg A$ enthält, sich $\neg A$ beweisen lässt. Ausgangspunkt muss aber ein direkt geführter Beweis für $A \Vdash (\neg A)$ bzw. $\Vdash A \Rightarrow \neg A$ sein.

Durch das Streichen der Regel (R4) entsteht der Kalkül **SKP⁻**. Er erzeugt eine echte Teilmenge der Theoreme von **SKP**. Die Theoreme der Korrektheit und der Nicht-Trivialität übertragen sich auf **SKP⁻**.

Da die parakonsistente Logik **SKP⁻** so schwach ist, dass direkte Beweise zu führen sind, heißt, dass die logische Kraft in mehr Prämissen liegen muss. Die betreffenden Theorien müssen eine größere Menge expliziter Annahmen machen, da sich z.B. Umkehrungen nicht mehr logisch ergeben¹³⁴.

Alternativ zum Preisgeben der Kontraposition könnte man behaupten (vgl. IC:109), dass der konditionale Junktor, der in Konvention (T) vorkommt, eben nicht das Enthaltensein ist. Und für diesen Junktor in Konvention (T) dann die Kontraposition nicht gilt. Dann müssten wir aber

¹³⁴ Priest behauptet gelegentlich zwar Reductio und Modus Tollens blieben quasi-gültig (vgl. "Reductio ad Absurdum et Modus Tollendo Ponens"). meint damit indessen, dass sie sofern sie in konsistenten Kontexten verwendet werden, gültig sind. Damit gehen

klären, welche Semantik dieses weitere Konditional besitzt und welche Ableitungsregeln es regieren und wie es mit dem Enthaltensein interagiert. Diese Lösung scheint einen außerordentlichen *ad hoc*-Charakter zu besitzen¹³⁵.

Aber es scheint noch schlimmer zu kommen.

Es drohen **Hyper-Widersprüche**. Hyper-Widersprüche sind - wenn sie sich zeigen ließen - das, was für die mehrwertige Logik die verstärkten Lügner waren: Antinomien. Und zwar solche, die sich nicht mit den Mitteln behandelt werden können, die für die ursprünglichen Antinomien entwickelt wurden. Ließen sich Hyperwidersprüche zeigen, wäre die parakonsistente Logik gescheitert. Wir müßten zur Hierarchie-Konzeption zurückkehren.

(Priest hat auch selbst auf dieses Problem aufmerksam gemacht.¹³⁶)

Das Problem ergibt sich durch die ursprüngliche Konstruktion, die Bewertungsfunktion bezüglich von Aussagen so aufzufassen, dass sie Mengen von Wahrheitswerten zuordnet (nämlich $\{0\}$, $\{1\}$ oder $\{0,1\}$), um dann zu sagen $1 \in v(A)$ oder $0 \in v(A)$. Die Erweiterung gegenüber der Standard-Semantik bestand darin, anstelle einer Funktion in $\{0,1\}$ Aussagen über der Potenzmenge von $\{0,1\}$ abzüglich der leeren Menge, $\emptyset$, zu bewerten¹³⁷. Neben die Stufe 0 der Standard-Aussagen tritt also eine Stufe 1

pragmatische Aspekte in die Logik ein, ohne dass die betreffenden Regeln für SKP zu retten wären, bei Beibehaltung der anderen Adäquatheitsbedingungen.

¹³⁵ Priest bezieht sich an dieser Stelle im Übrigen auf Stalnakers konditionalen Junktor für kontrafaktische Konditionale. Damit verbindet sich aber nicht nur eine Ähnlichkeitsrelation zwischen möglichen Welten, die in Semantik II bis jetzt nicht vorliegt, sondern auch *weitere* Beschränkungen der Logik. So gilt für Stalnakers Konditional *keine Transitivität*, die für das Enthaltensein gilt (vgl. Stalnaker, "A Theory of Conditionals").

¹³⁶ vgl. zum folgenden: Priest, "Can Contradictions Be True?"; ders. "Hyper-Contradictions"; Everett, "A Note on Priest's 'Hyper-Contradictions'"; Priest, "Everett's Trilogy"; Smiley, "Can Contradictions Be True?".

¹³⁷ Die leere Menge fällt heraus, da in der parakonsistenten Semantik bei Priest keine Wahrheitswertlücken angenommen werden. Eine parakonsistente Semantik mit Wahrheitswertlücken würde auch $\emptyset$ als Bewertung von Aussagen zulassen. Diese Problematik ist aber unabhängig vom eigentlichen Anliegen der parakonsistenten Logik.

der Antinomien, die in Stufe 0 auftreten und die den Wert $\{0,1\}$ erhalten. Daneben eine Stufe 2 der Antinomien, die auf Stufe 1 auftreten, wie

(7) Aussage (7) ist nur falsch.

Betrachten wir diese Aussagen nun nämlich von der Semantik her - und nicht wie oben bezüglich (2) von der syntaktischen Seite, ob er die Form einer Dialetheia mit sich bringt. (7) ist wahr oder falsch, da das Tertium Non Datur gilt. Wenn (7) wahr ist, ist die zugeordnete Menge von Wahrheitswerten entweder $\{1\}$ oder $\{0,1\}$ aufgrund der Definition von "() ist wahr". Hinsichtlich (7) besagt die Wahrheit von (7) aber, dass das, was (7) behauptet, der Fall ist, d.h. (7) ist nur falsch, d.h. die zugeordnete Menge von Wahrheitswerten ist $\{0\}$. Entweder ist also $\{1\}$ mit $\{0\}$ identisch oder $\{0,1\}$ ist mit $\{0\}$ identisch, und das heißt in beiden Fällen $1=0$, was nicht nur absurd ist, sondern die Trivialität der Sprache mit sich bringt, da nun jede Aussage zugleich wahr und falsch wäre. Ist (7) hingegen nur falsch, d.h. $v(7)=\{0\}$, dann müsste aufgrund der Falschheit der Behauptung von (7) die Bewertung von (7) $\{1\}$ oder $\{0,1\}$ sein, was wieder dazu führt, dass $1=0$. Eine Antinomie - und zwar eine Antinomie einer neuen Art!

Der einzige Ausweg scheint darin zu liegen, dass (7) zugleich $\{0\}$ und $\{1\}$ als Bewertungen hat. Das heißt dass $v(7)=\{\{0\},\{1\}\}$. Eine solche Bewertung kommt auf der Stufe 1 aber nicht vor. Sie stellt einen für Aussagen der Stufe 1 unmöglichen Wahrheitswert da, da diese Menge nicht in der Potenzmenge der Wahrheitswerte der Stufe 0 vorkommt. Es handelt sich aber um einen möglichen Wahrheitswert auf Stufe 2, da auf Stufe 2 die Potenzmenge der Wahrheitswertmenge von Stufe 1 zur Verfügung steht, d.h. die Potenzmenge von $\{\{1\},\{0\},\{0,1\}\}$. Und $\{\{0\},\{1\}\}$ kommt in dieser Potenzmenge vor. (7) zwingt uns also - anlog dem Lügner - eine Stufe in der Bewertung aufzusteigen. Die syntaktische Behandlung - wie wir sie bei Aussage (3) betrachtet haben - zeigt zwar zu Recht, dass es sich hier um eine Dialetheia handelt, doch muss die entsprechende Bewertung eine Stufe höher vorgenommen werden. Und da liegt die Problematik.

Denn nun lässt sich das Verfahren wiederholen - *und wir erhalten eine Bewertungshierarchie*¹³⁸!

Das bringt alle Probleme von semantischen Hierarchien zurück in die parakonsistente Logik: "Wahrheit" wird also stufenabhängig verwendet? Von welcher Stufe reden wir eigentlich über die Gesamtkonstruktion? - usw.

Priest versucht dann zu zeigen, dass bezüglich der Folgerungsrelation nach der Änderung, die zwischen Stufe 0 und Stufe 1 eintritt, nichts Neues mehr passiert¹³⁹. Doch die dazu nötige Konstruktion lässt wieder eine Aussage zu, die sich als Hyper-Widerspruch verhält¹⁴⁰.

Deshalb muss das Verständnis der Interpretations- bzw. Bewertungsfunktion grundsätzlich geändert werden. Die Bewertung bzw. Interpretation muss von einer *Relation* vorgenommen werden¹⁴¹. Dabei ist von der Menge der Wahrheitswerte $\{0,1\}$ auszugehen, wie es die Standard-Semantik auch tut. Im Unterschied zu dieser faßt die parakonsistente Semantik die Bewertung nun als *Relation* auf, das heißt eine Aussage kann auch zu beiden Elementen der Menge der Wahrheitswerte in Beziehung gesetzt werden. Als relationale Schreibweise bietet sich an: $v(A,1)$, d.h. die Bewertung v ordnet A den Wert 1 zu. Dabei gilt insbesondere

$$(8) \neg (v(A,1) \Rightarrow \neg v(A,0)) \wedge \neg (v(A,0) \Rightarrow \neg v(A,1))$$

d.h. die Bewertung einer Aussage mit "1" schließt nicht ihre Bewertung mit "0" aus, und umgekehrt. Wir können definieren:

$$(D1) \text{ A ist wahr in } M_i := v_i(A,1)$$

$$(D2) \text{ A ist gültig} := (\forall v)v(A,1)$$

Bezogen auf diese Semantik hätte eine der Aussage (7) analoge verstärkte Lügner-Aussage die Form

$$(9) v((9),1) \Leftrightarrow (v((9),0) \wedge \neg v((9),1))$$

¹³⁸ So in Priest, "Hyper-Contradictions", S.238ff.

¹³⁹ vgl. ebd. S.240f.

¹⁴⁰ vgl. Everett, "A Note on Priest's 'Hypercontradiction'".

¹⁴¹ vgl. Priest, "Can Contradiction Be True?", S.50f.

Und diese Aussage (9) kommt schlimmstenfalls wieder

$$(10) v((9),1) \wedge v((9),0)$$

gleich, wobei *jetzt* die beiden Wahrheitswerte der Aussage (9) auf der einen semantischen Stufe, die wir hier haben, zugewiesen werden.

Damit ist die Drohung der Hyper-Widersprüche im Ansatz zurückgewiesen.¹⁴²

Alles was in vorherigen Abschnitten über Bewertungen von Aussagen gesagt wurde, muss also so verstanden werden, dass eigentlich von Relationen gesprochen wird.¹⁴³

¹⁴² "im Ansatz" meint, dass die große Drohung abgewendet wurde und nun einzelne Vorschläge für Hyper-Widersprüche zu widerlegen sind. Beispielsweise lässt sich nun nicht einfach folgende neue Art von Hyper-Widersprüchen einführen (vgl. Priest, "Can Contradictions Be True?", S.51): Wir können zwar jetzt die Menge d der Wahrheitswerte definieren, die eine Aussage A erhält, nämlich $d=\{x|v(A,x)\}$. Doch wäre nun für einen Hyper-Widerspruch zu zeigen, dass wir zugleich $1 \in d$ und $d=\{0\}$ haben könnten. Letzteres käme $x=0 \Leftrightarrow v(A,x)$ gleich. Selbst wenn $1 \in d$ für einen verstärkten Lügner α gilt, kommen wir nicht zu $d=\{0\}$ bzw. der zweiten Aussage. Denn hier müßte man argumentieren:

$$\neg v(\alpha,1), \text{ also } x \neq 1, \text{ also } x=0, \text{ da } v(\alpha,x) \Rightarrow (x=0 \vee x=1)$$

wobei dieses Argument aber auf dem parakonsistent ungültigen disjunktiven Syllogismus beruht. Everett hat in einer unpublizierten Arbeit, von der Priest in "Everett's Trilogy" berichtet (ebd.S.639f.), versucht einen anderen Hyper-Widerspruch für die relationale Auffassung der Bewertung zu konstruieren, aber auch dort wird von einer Schlußregel Gebrauch gemacht (der Exportation über das Enthaltensein, $A \wedge B \Rightarrow C \models A \Rightarrow (B \Rightarrow C)$), die parakonsistent ungültig ist, wie man an dem folgenden **SKP** Semantik II Modell III sehen kann:

$$\begin{array}{lcl} @ & = & = > w1 \\ A \wedge B & & B \\ C & & \end{array}$$

Die Prämisse der Exportation über das Enthaltensein ist in @ wahr, da in allen Welten, in denen $A \wedge B$ wahr ist, auch C wahr ist, aber die Konsequenz gilt nicht, da in @, obwohl A gilt,

$B \Rightarrow C$ nicht gilt, da @ $w1$ sehen kann, d.h. eine B -Welt, die keine C -Welt ist. Damit ist diese Regel nach (SIF) ungültig.

¹⁴³ Aus Gründen der Darstellung habe ich mit den üblichen Bewertungsfunktionen angefangen. Die tatsächliche Neuformulierung der bisherigen Semantiken können wir uns hier ersparen, wenn wir festhalten, dass sie eigentlich Relationen benutzen, um Hyperwidersprüche zu vermeiden. Die Änderung in der Definition von Gültigkeit etc. sind marginal.

Weitere stark anti-triviale Kalküle

Es gibt weitere Vorschläge für stark antitriviale parakonsistente Kalküle.

- Beispielsweise stellen Arruda und da Costa einige entsprechende Relevanzkalküle vor¹⁴⁴.

Das System **P** besteht aus $A \supset A$, dem Tertium Non Datur, den Axiomen der Konjunktionsbeseitigung, der Disjunktionseinführung und der Negationsbeseitigung als Axiomen sowie den Regeln des Modus Ponens, der Transitivität, der Konjunktionseinführung und

$$A \supset B, A \supset C \vdash (A \supset B \wedge C)$$

$$A \wedge B \supset D, A \wedge C \supset D \vdash (A \wedge (B \vee C) \supset D).$$

P* ergibt sich durch die Ergänzung um die *reductio*:

$$A \supset B, A \supset \neg B \vdash (\neg A).$$

In **P*** gilt dann der Satz vom Widerspruch. Viele Standard-Theoreme gelten nicht: Doppelte Negation, Transitivität als Theorem, Adjunktion als Theorem, Kontraposition, Modus Ponens-Theorem, Importation, Exportation, hypothetische Abschwächung, *reductio* als Theorem - usw.¹⁴⁵

Schwerwiegender ist, dass bezüglich der Junktoren eine 5-wertige Semantik angegeben wird. So wird z.B. definiert¹⁴⁶:

¹⁴⁴ vgl. zum folgenden: Arruda/da Costa, "On the Relevant Systems P and P* and Some Related Systems".

¹⁴⁵ Einige von den Theoremen, die in **P*** nicht gelten, gelten in **SKP**.

¹⁴⁶ Um Platz zu sparen, wähle ich hier eine andere Darstellung von Wahrheitstafeln: die oberste Reihe faßt die Bewertungsmöglichkeiten von B, die erste Spalte die Bewertungsmöglichkeiten von A, die Eintragungen in der Matrix ergeben sich durch die jeweiligen Kombinationen.

$A \wedge B$		1	2	3	4	5
1		1	1	4	4	1
2		1	2	4	4	2
3		4	4	3	4	3
4		4	4	4	4	4
5		1	2	3	4	5

Was soll uns das sagen? Irgendeine Intuition bezüglich der Konjunktion finde ich darin nicht wieder. Die Semantik wird auch nicht weiter begründet. Darf man aber nicht um eines Vollständigkeitsbeweises willen so vorgehen? Nein: Wenn wir die deduktive Vollständigkeit eines Kalküls zeigen wollen, müssen wir zeigen, dass er alle logischen Wahrheiten, die sich mit seinem Vokabular erzeugen lassen, ableiten kann; wenn eine Aussage logisch wahr ist, dann ist sie wahr unabhängig von der Interpretation, doch festgehalten werden muss an der Interpretation der logischen Ausdrücke, diese bestimmen ja die Semantik, relativ zu der deduktive Vollständigkeit gezeigt werden soll.

Diese Arruda/da Costa-Semantik gibt daher ein Musterbeispiel einer philosophisch irrelevanten Semantik ab. Arruda und da Costa können zwar beweisen, dass diese Systeme nicht "finit trivialisierbar" sind (d.h. es keine Aussage gibt, aus der beliebige Aussagen folgen), doch hat dies bestenfalls metalogische Signifikanz im Sinne der Frage, was für Systeme mit was für Semantiken man überhaupt definieren und bezüglich ihres Verhaltens untersuchen kann.

- Das System **H** von Routley¹⁴⁷ enthält anders als **P** Kontraposition als Regel und die DeMorgan Gesetze. Die Kontraposition haben wir indes- sen gerade in Kapitel 4.3 als äußerst problematisch kennengelernt.

¹⁴⁷ vgl. Routley/Loparic, "Semantical Analysis of Arruda-da Costa **P** Systems and Adjacent Non-Replacement Relevant Systems".

Es gelten in **H** laut Arruda und da Costa nicht nur das Tertium Non Datur und Negationsbeseitigung nicht, vielmehr *ist überhaupt keine Aussage der Form $\neg A$ gültig!*¹⁴⁸

Das ist ebenfalls philosophisch inakzeptabel. Weder würde der Satz vom Widerspruch noch irgendeine andere Negation einer logischen Falschheit gelten.

¹⁴⁸ Arruda/da Costa, "On the Relevant Systems P and P* and Some Related Systems", S.43. In ihrem Beweis verwenden sie allerdings eine Wahrheitstafel für die Negation, in der die Negation immer (!) falsch ist. Da sie zugleich die Standard-Tafel der Implikation verwenden (!), muss bei Vollständigkeit des Systems *ex contradictione quod libet* als Theorem gelten und für die Anti-Trivialität müsste der Modus Ponens aufgegeben werden! Für den Nachweis, dass negationsfreie Theoreme, die in **P*** nicht gelten, auch in **H** nicht gelten, *wechseln sie dann die Semantik!*