

BAYERISCHE AKADEMIE DER WISSENSCHAFTEN
MATHEMATISCH-NATURWISSENSCHAFTLICHE KLASSE

SITZUNGSBERICHTE

JAHRGANG

1959

MÜNCHEN 1960

VERLAG DER BAYERISCHEN AKADEMIE DER WISSENSCHAFTEN

In Kommission bei der C. H. Beck'schen Verlagsbuchhandlung München

Untersuchungen zur Quantenlogik

Von Peter Mittelstaedt in München

Vorgelegt von Herrn Werner Heisenberg, am 9. Oktober 1959

Übersicht

I. Einleitung und Problemstellung	322
II. Die Logik der kommensurablen Eigenschaften	326
1 Problemstellung	326
2 Die logischen Begriffe	328
a) Die Halbordnung 328 – b) Die Konjunktion 329 – c) Die Disjunktion 330 – d) Die Implikation 331 – e) Die Negation 333	
3 Der Logikkalkül	334
III. Die Logik der inkommensurablen Eigenschaften	338
1 Physik und Logik	338
2 Die logischen Begriffe.	343
a) Die Halbordnung 345 – b) Die Konjunktion 347 – c) Die Disjunktion 348 – d) Die Implikation 350 – e) Die Negation 353	
f) Quantoren 354	
3 Der Aufbau der Quantenlogik	354
a) Die effektive Quantenlogik 354 – b) Distributivität und Modularität 356 – c) Negation und Komplement 360 – d) Diskussion einiger Beispiele 362	
IV. Der quantenlogische Modalkalkül	364
1 Sprache und Metasprache	364
2 Der Aufbau des Modalkalküls	367
3 Möglichkeit und Wirklichkeit	370
Anhang I. Quantentheorie.	374
1 Abgeschlossene Linearmannigfaltigkeiten	374
2 Projektionsoperatoren	375
3 Quantenmechanische Eigenschaften	377
Anhang II. Verbandstheoretische Hilfsmittel	379
Anhang III. Operative Logik	382
Literatur	385

I. Einleitung und Problemstellung

Um die Fragestellung der vorliegenden Untersuchung zu verdeutlichen, ist es zweckmäßig, kurz auf die Entwicklung des Problems der Quantenlogik einzugehen. Im Anschluß an die durch J. v. Neumann [1] angegebene Formulierung der Quantenmechanik (vgl. Anh. I) wurde im Jahre 1936 von Birkhoff und v. Neumann [2] eine Theorie der abgeschlossenen Linear-mannigfaltigkeiten des Hilbertraumes entwickelt, welche von einem algebraisch-verbandstheoretischen Gesichtspunkt her gesehen (vgl. Anh. II) in vieler Hinsicht eine gewisse Ähnlichkeit mit der sog. „klassischen“ Logik besitzt. Aus diesem Grunde, und wegen der physikalischen Deutung, welche abgeschlossene Linear-mannigfaltigkeiten im Rahmen der Quantentheorie erfahren [1], identifizierten Birkhoff und v. Neumann diese Unterräume mit Eigenschaften (qualities) eines physikalischen Systems, so daß der oben erwähnte Formalismus, die sog. Quantenlogik, als eine Logik quantenmechanischer Eigenschaften eines Systems interpretiert werden konnte. Die Abweichungen dieser Quantenlogik von der klassischen Logik sind derart, daß sich in ihnen alle wesentlichen, im Vergleich mit der klassischen Physik ursprünglich als paradox empfundenen Eigenschaften der Quantenmechanik in einer sehr allgemeinen Form darstellen lassen, woraus sich die grundlegende Bedeutung der Quantenlogik für alle Interpretationsfragen der Quantentheorie erklärt.

Im Rahmen der von Hilbert u. a. [3], [4] ausgearbeiteten axiomatischen Auffassung der Logik ließ sich diese Quantenlogik verstehen als eine kontingente Theorie, die für einen gewissen Typus von Aussagen, eben die quantentheoretischen Eigenschaften, Gültigkeit besitzt und daher die für diesen Gegenstandsbereich zutreffende Logik darstellt.

Während sich so die Quantenlogik in die von der axiomatischen Methode bestimmte Vorstellung über das Wesen der Logik widerstandslos einordnen läßt, ist dies bei Berücksichtigung der neueren Entwicklung der Logik nicht mehr der Fall. Die im Anschluß

an Skolem [5], Curry [6], insbesondere von Lorenzen [7], [8], [9] entwickelte operative Auffassung der Logik begreift die Gesetze der Logik nicht mehr als willkürlich gewählte, lediglich dem jeweiligen Sachgebiet angepaßte Behauptungen, sondern als Regeln, deren Evidenz sich jetzt aus einer Reflexion über die Möglichkeiten des schematischen Operierens mit Figuren nach gewissen Regeln ergibt (vgl. Anh. III). Ein solches System von Regeln zum Operieren mit Figuren wird dabei als ein Kalkül bezeichnet. Dadurch verliert der logische Formalismus jegliche Willkür, da von jedem behaupteten logischen Satz jetzt zu zeigen ist, wie er sich operativ begründen läßt. An Stelle der erwähnten Kontingenz der logischen Sätze tritt somit eine durch Reflexion über das Operieren in Kalkülen gewonnene Evidenz dieser Behauptungen, gleichzeitig jedoch wird klar, daß diese operative Logik entsprechend ihrer Begründung zunächst nur auf Kalküle anwendbar ist. Jedoch ist damit die Verwendung der operativen Logik in der gesamten Mathematik legitimiert.¹

Von hier aus wird die Problemstellung der vorliegenden Arbeit sichtbar. Es erhebt sich nämlich zunächst die Frage, ob der von Birkhoff und v. Neumann entwickelte Formalismus in diesem oder einem verallgemeinerten operativen Sinne als eine Logik bezeichnet werden kann. Die Behauptung, daß in dem von der Quantentheorie beschriebenen Bereich der physikalischen Wirklichkeit eine andere Logik, die Quantenlogik, gilt, wird aber dadurch etwas problematisch, daß die Quantenmechanik eine mathematisch streng formulierbare Theorie ist, für die, als mathematischen Formalismus, die operative Logik gültig sein soll. Es stellt sich somit die weitere Frage, wieso die Logik der quantentheoretischen Aussagen, falls sie sich als eine Logik im operativen Sinne auffassen läßt, in einigen wesentlichen Punkten von der operativen Logik der Kalküle verschieden ist.

Um diese Fragen zu entscheiden, werden wir umgekehrt vorgehen und zunächst das Problem untersuchen, inwieweit Aus-

¹ Die operative Logik (Anh. III) umfaßt die effektive Logik, nicht hingegen das tertium non datur, das nur für spezielle Kalküle als zulässig nachgewiesen werden kann.

sagen¹ über die physikalische Wirklichkeit den Gesetzen der Logik gehorchen. Unter Logik wollen wir dabei stets die effektive bzw. die daraus durch Hinzunahme des tertium non datur entstehende fiktive oder klassische Logik verstehen. Unser Wissen über die Wirklichkeit ist, jedenfalls für einen sehr weiten Bereich physikalischer Phänomene, in der Quantentheorie und der dazu gehörigen Theorie des Meßprozesses [10] zusammengefaßt. Wir werden daher die Frage zu untersuchen haben, ob die durch diese Theorien charakterisierten Aussagen über die Natur sich auf konstruierbare Figuren (Aussagen) eines Kalküls zurückführen lassen. Wir werden sehen, daß dies nicht allgemein der Fall ist. Vielmehr werden die operativen Möglichkeiten, die sonst in Kalkülen bestehen, durch den Bezug der Aussagen auf die physikalische Wirklichkeit einigen wesentlichen Einschränkungen unterworfen, so daß einige der Gesetze, die in der effektiven Logik gültig sind, für den Bereich physikalischer Aussagen nicht mehr allgemein zutreffen. Diese so begründete Logik werden wir als operative oder effektive Quantenlogik bezeichnen.

Es gibt allerdings eine bestimmte Klasse von physikalischen Aussagen, bei denen die Reduktion auf die Theorie der Kalküle doch durchgeführt werden kann, die im Kap. II behandelten kommensurablen Eigenschaften, deren Logik daher auch mit der klassischen Logik genau übereinstimmt. Dieser Spezialfall ist deshalb von großer Bedeutung, weil sich die gesamte makroskopische Physik auf die Theorie der kommensurablen Eigenschaften zurückführen läßt, so daß die im Kap. II durchgeführte Diskussion u. a. zeigt, daß in der klassischen Physik stets die klassische Logik gilt. Gleichzeitig soll Kap. II dazu dienen, die Begriffsbildungen der Quantenlogik und der operativen Logik an einem besonders einfachen Spezialfall einzuführen.

Zieht man über die operative Quantenlogik hinaus noch die für quantenmechanische Aussagen stets bekannte explizite Ge-

¹ Was dabei im einzelnen unter „Aussage“ zu verstehen ist, soll weiter unten genauer präzisiert werden. Wir wollen uns hier auf die vorläufige Formulierung beschränken, daß „Aussagen“ sich auf die Ergebnisse irgendwelcher physikalischer Experimente beziehen sollen.

stalt dieser Aussagen in Betracht, so lassen sich zunächst einige weitere Sätze beweisen, wodurch wir zur eigentlichen Quantenlogik geführt werden. Das tertium non datur ist von dieser Art. Von hier aus ist es dann auch möglich, den Zusammenhang quantenmechanischer Aussagen mit den von Birkhoff und v. Neumann [2] verwendeten abgeschlossenen Linearmannigfaltigkeiten herzustellen. Es wird sich zeigen, daß man auf diesem Wege exakt die von Birkhoff und v. Neumann angegebene Quantenlogik erhält. Insbesondere ist bemerkenswert, daß die in dieser Theorie zunächst auf Grund rein mathematischer Überlegungen an den abgeschlossenen Linearmannigfaltigkeiten eingeführten logischen Begriffe „und“, „oder“ und „nicht“ genau den in der operativen Logik eingeführten Sinn haben. Man wird daher die Quantenlogik als eine Logik im operativen Sinne des Wortes ansehen dürfen, wobei allerdings der Unterschied zur klassischen Logik dadurch entsteht, daß die operativen Möglichkeiten gewissen, durch den Bezug auf die physikalische Wirklichkeit bedingten Einschränkungen unterworfen sind. In diesem Sinne wird man die vorliegende Untersuchung als eine operative Begründung der Quantenlogik auffassen dürfen.

In einem abschließenden Kapitel (Kap. IV) soll, ausgehend von der in Kap. II, III untersuchten Quantenlogik, eine auf dieser für das inhaltliche Verständnis etwas sehr abstrakten Objektsprache aufgebaute Metasprache eingeführt werden, die wohl am ehesten derjenigen Sprache entspricht, in der üblicherweise die Phänomene der Quantentheorie beschrieben werden. Auf diese sehr wesentliche Eigenschaft der Metasprache hat insbesondere v. Weizsäcker hingewiesen [11], [12]. Die Einführung des für die Quantentheorie entscheidenden Begriffes der Möglichkeit in die Metasprache führt dann zwangsläufig zu einem auf der Quantenlogik aufbauenden Modalkalkül.¹

¹ Die mit der mehrwertigen Logik [13], [14], [15] zusammenhängenden Fragen werden in einer späteren Arbeit behandelt. (Erscheint im Archiv für mathematische Logik und Grundlagenforschung).

II. Die Logik der kommensurablen Eigenschaften

II. 1 Problemstellung

In dem vorliegenden Kapitel soll die Logik derjenigen Aussagen A, B, C untersucht werden, die besagen, daß ein im Zustand $|f\rangle$ befindliches quantenmechanisches System S bestimmte Eigenschaften E_A, E_B, E_C besitzt, wobei angenommen werden soll, daß das Vorliegen dieser Eigenschaften gleichzeitig an dem betreffenden System durch eine Messung nachgewiesen werden kann. Unter einer Aussage A soll daher (vgl. Anhang I) die Behauptung verstanden werden, daß der Zustand $|f\rangle$ des betrachteten Systems in einer abgeschlossenen Linearmannigfaltigkeit M_A des Hilbertraumes liegt. Ist P_A der Projektionsoperator, der auf M_A projiziert, so werde also A durch

$$A \rightleftharpoons P_A |f\rangle = |f\rangle$$

definiert. (Das Symbol $\rightleftharpoons$ verwenden wir stets für Definitionen).

Die Logik dieser A, B, C soll im folgenden kurz als die Logik der kommensurablen Eigenschaften bezeichnet werden. Ebenso werden wir der Einfachheit halber meist von den „Eigenschaften“ A, B, C sprechen, soweit es keine Verwechslungsmöglichkeiten mit den oben definierten E_A geben wird.

Die Messung einer beliebigen Eigenschaft E_A an einem System S verwandelt im allgemeinen dessen Zustand $|f\rangle$ in einen prinzipiell nicht näher bestimmten Eigenzustand $|A_i\rangle$ von P_A . Der statistische Operator $P = |f\rangle \langle f|$ von S geht nämlich durch die Wechselwirkung des Systems mit dem Meßgerät in den Operator $\sum_i |A_i\rangle \langle A_i| |\langle A_i | f \rangle|^2$ eines Gemenges über[10]. Es ist daher nicht möglich, das Vorliegen mehrerer allgemeiner Eigenschaften E_A, E_B, E_C gleichzeitig an einem Zustand $|f\rangle$ durch eine Messung nachzuweisen. Die hier vorgenommene Beschränkung auf kommensurable Eigenschaften bedeutet, daß die Messung dieser Größen den Zustand $|f\rangle$ nicht verändert, so daß mehrere solche Eigenschaften simultan an S gemessen werden können. Mathematisch heißt das, daß die entsprechenden Projektionsoperatoren unter sich und mit P paarweise vertauschbar sind.

Der Grund dafür, daß hier zunächst der Spezialfall der kommensurablen Eigenschaften besprochen werden soll, ist, daß sehr viele der für die eigentliche Quantenlogik wichtigen Begriffsbildungen bereits an diesem übersichtlichen Fall diskutiert werden können.

Da alle Projektionsoperatoren $P_A, P_B \dots$ mit P vertauschbar sind, folgt zunächst, daß das Vorliegen einer Eigenschaft E_A an einem System genau dann experimentell verifiziert werden kann, wenn $P_A |f\rangle = |f\rangle$ gilt. Die in der Quantentheorie ableitbaren Aussagen $P_A |f\rangle = |f\rangle$ geben daher eine vollständige Information über den Ausgang von Experimenten, die zur Feststellung kommensurabler Eigenschaften dienen. Die oben definierten Aussagen A, B, C entsprechen also eindeutig den ableitbaren Aussagen des quantentheoretischen Kalküls. Die Logik dieser A, B, C wird somit die Logik der Aussagen eines beliebigen Kalküls sein, d. h. die effektive Logik. Wir können daher den systematischen Aufbau der Logik kommensurabler Eigenschaften dazu verwenden, die Methoden und Begriffsbildungen der operativen Logik an einem für uns besonders interessanten Beispiel zu demonstrieren. Aus diesem Grunde soll der hier vorgenommene Aufbau dieser Logik auch etwas ausführlicher erfolgen, als dies der Sache nach unbedingt erforderlich wäre.

Zur Untersuchung der algebraischen Struktur der im folgenden diskutierten logischen Systeme werden wir uns weitgehend verbandstheoretischer Mittel bedienen (vgl. Anh. II). Da der in diesem Kapitel untersuchte Kalkül auch in dieser Hinsicht einen besonders einfachen Spezialfall darstellt, sollen die in dem folgenden Kapitel benutzten verbandstheoretischen Begriffe an diesem Beispiele genauer erklärt werden.

Der so hergeleitete Kalkül der kommensurablen Eigenschaften wird genau mit der effektiven Logik irgendwelcher Kalkülausagen übereinstimmen. Darüber hinaus lassen sich jedoch noch einige zusätzliche Regeln beweisen, da man hier außerdem die explizite Form der Aussagen kennt. Das tertium non datur, das effektiv nicht beweisbar ist, gehört zu diesen Regeln. Es sollen jedoch zuerst die effektiv beweisbaren Sätze angegeben werden, und die explizite Gestalt der Aussagen nur zum Beweis der zu-

sätzlichen Regeln verwendet werden. Lediglich in einigen Fällen, in denen die Beweise der effektiv gültigen Gesetze durch Benutzung der expliziten Gestalt der Aussagen besonders durchsichtig werden, sollen zusätzlich auch noch diese Regeln angegeben werden.

Unter Benutzung der expliziten Gestalt wird es weiter möglich sein, die mit den logischen Partikeln zusammengesetzten Aussagen $A \wedge B$, $A \vee B$, $A \rightarrow B$ explizit anzugeben. Es wird nämlich möglich sein, mit Hilfe der Operatoren P_A und P_B sehr einfache zusammengesetzte Operatoren zu definieren, die den zusammengesetzten Aussagen entsprechen. Von hier aus wird man dann einen einfachen Zugang zu der Möglichkeit finden, zusammengesetzten Aussagen gewisse abgeschlossene Linearmannigfaltigkeiten des Hilbertraumes zuzuordnen, ein Verfahren, von dem insbesondere bei der Behandlung der eigentlichen Quantenlogik in Kap. III ausführlich Gebrauch gemacht werden soll, da man auf diese Weise einen unmittelbaren Zugang zu der Logik von Birkhoff und v. Neumann [2] gewinnt.

Um die aus den verschiedenen als gültig erwiesenen logischen Regeln sich ergebende algebraische Struktur besser charakterisieren zu können, werden wir uns, ebenso wie bei der effektiven Logik (Anh. III) verbandstheoretischer Begriffe bedienen. Es sei jedoch betont, daß die Verbandstheorie auch hier immer nur zur Illustration der gewonnenen Zusammenhänge verwendet wird, niemals aber zu deren Begründung.

II. 2 Die logischen Begriffe

a) Die Halbordnung

Wir gehen aus von Aussagen der Form

$$(II, 1) \quad A \rightleftharpoons P_A |f\rangle = |f\rangle$$

Zwischen diesen Aussagen möge eine Relation $A \rightarrow B$ definiert werden, die genau den in Anh. III eingeführten operativen Sinn hat. In bezug auf die physikalische Bedeutung von A und B heißt „ $\rightarrow$ “, daß wenn eine Messung A ergeben hat, dann auch B bei

einer geeigneten Messung als zutreffend vorgefunden wird. Bezüglich der Relation $A \rightarrow B$ gelten dann die Regeln

$$(II, 2) \quad A \rightarrow B$$

$$(II, 3) \quad A \rightarrow B, \quad B \rightarrow C \Rightarrow A \rightarrow C$$

Zum Beweis hat man zu zeigen, daß die Benutzung dieser Regeln aus jeder Ableitung eliminiert werden kann (vgl. Anh. III). Wir wollen hierauf nicht näher eingehen. Benutzt man die explizite Form (II, 1) der Aussagen, so ist (vgl. A I) $A \rightarrow B$ notwendig und hinreichend dafür, daß

$$(II, 4) \quad P_A = P_A P_B = P_B P_A$$

gilt. Unter Benutzung von (II, 4) läßt sich (II, 2, 3) sehr einfach beweisen: (II, 2): $P_A = P_A^2$; (II, 3): $P_A = P_A P_B$, $P_B = P_A P_C$, also $P_A = P_A P_B P_C = P_A P_C$.

Bezüglich der abgeschlossenen Linearmannigfaltigkeiten ist (II, 4) äquivalent zu

$$(II, 5) \quad M_A \leq M_B$$

Vom verbandstheoretischen Gesichtspunkt bilden die Aussagen A, B, C bezüglich der Relation „ $\rightarrow$ “ wegen (II, 2, 3) eine Halbordnung.

b) Die Konjunktion

Sind A, B Aussagen im obigen Sinne, so führen wir die Konjunktion durch die Regel

$$(II, 6) \quad A, B \rightarrow A \wedge B$$

ein. In dem Kalkül der Aussagen A, B , der nicht das Zeichen $\wedge$ enthält, ist dann (II, 6) relativ zulässig. Durch Inversion gewinnt man die weiteren Regeln

$$(II, 7) \quad A \wedge B \rightarrow A$$

$$(II, 8) \quad A \wedge B \rightarrow B$$

$$(II, 9) \quad C \rightarrow A, \quad C \rightarrow B \Rightarrow C \rightarrow A \wedge B$$

Umgangssprachlich würde man $A \wedge B$ mit „ A und B “ übersetzen. Unabhängig von (II, 6) können wir (II, 7, 8, 9) als Definition der Aussage $A \wedge B$ auffassen. Unter Verwendung der expliziten Form (II, 1) folgt aus (II, 7, 8, 9)

$$(II, 7a) \quad P_{A \wedge B} = P_A \cdot P_{A \wedge B}$$

$$(II, 8a) \quad P_{A \wedge B} = P_B \cdot P_{A \wedge B}$$

$$(II, 9a) \quad P_A P_C = P_C, \quad P_B P_C = P_C \Rightarrow P_{A \wedge B} P_C = P_C$$

Andererseits gilt für den Operator $P_A P_B$

$$(II, 7b) \quad P_A P_B = P_A (P_A P_B)$$

$$(II, 8b) \quad P_A P_B = P_A (P_A P_B)$$

$$(II, 9b) \quad P_A P_C = P_C, \quad P_B P_C = P_C \Rightarrow P_A P_B P_C = P_C$$

so daß aus (II, 9a, 9b)

$$(II, 10) \quad P_{A \wedge B} = P_A \cdot P_B$$

folgt. Nach Anh. I projiziert $P_A \cdot P_B$ auf den Durchschnitt $M_A \cap M_B$ der zu P_A und P_B gehörigen Unterräume M_A und M_B . Vom verbandstheoretischen Gesichtspunkt gesehen bilden die Aussagen bezgl. $\wedge$ einen Halbverband wegen (II, 7, 8, 9).

c) Die Disjunktion

Durch die Regeln

$$(II, 11) \quad A \rightarrow A \vee B$$

$$(II, 12) \quad B \rightarrow A \vee B$$

werden die Aussagen $A \vee B$ eingeführt. (II, 11, 12) sind relativ zulässig in einem Kalkül, der nicht das Zeichen $\vee$ enthält. Durch Inversion gewinnt man die Regel:

$$(II, 13) \quad A \rightarrow C, \quad B \rightarrow C \Rightarrow A \vee B \rightarrow C$$

Umgangssprachlich würde man $A \vee B$ mit „ A oder B “ übersetzen. (II, 11, 12, 13) kann jetzt wieder als Definition der Aussage

$A \vee B$ aufgefaßt werden, denn unter Benutzung der expliziten Form (II, 1) der Aussagen folgt, daß

$$(II, 11a) \quad P_A = P_A \cdot P_{A \vee B}$$

$$(II, 12a) \quad P_B = P_B \cdot P_{A \vee B}$$

$$(II, 13a) \quad P_A = P_A P_C, P_B = P_B P_C \Rightarrow P_{A \vee B} = P_{A \vee B} P_C$$

gilt. Weiterhin gilt, daß

$$(II, 11b) \quad P_A = P_A (P_A + P_B - P_A P_B)$$

$$(II, 12b) \quad P_B = P_B (P_A + P_B - P_A P_B)$$

$$(II, 13b) \quad P_A = P_A P_C, P_B = P_B P_C \Rightarrow P_A + P_B - P_A P_B = \\ = (P_A + P_B - P_A P_B) P_C$$

so daß aus (II, 13a, 13b) folgt, daß

$$(II, 14) \quad P_{A \vee B} = P_A + P_B - P_A P_B$$

gilt. $P_{A \vee B}$ ist damit eindeutig festgelegt.

Aus Anh. I folgt weiter, daß $P_{A \vee B}$ auf den von M_A und M_B aufgespannten Teilraum $M_A \cup M_B$ projiziert. Vom verbandstheoretischen Gesichtspunkt bilden die Aussagen $A, B, C \dots$ bzgl. $\wedge$ und $\vee$ einen Verband, oder bzgl. $\rightarrow$ eine Halbordnung, in der für je zwei Elemente A und B eine obere Grenze $A \vee B$ und eine untere Grenze $A \wedge B$ existiert.

d) Die Implikation

Für die oben behandelte Implikation „ $\rightarrow$ “ gelten außer (II, 2, 3) in der operativen Logik noch die beiden weiteren Relationen

$$(II, 15) \quad A \dot{\wedge} A \rightarrow B \dot{\rightarrow} B$$

$$(II, 16) \quad A \wedge C \rightarrow B \Rightarrow C \dot{\rightarrow} A \rightarrow B$$

Beide Gesetze können wiederum dadurch als gültig bewiesen werden, daß sich ihre Benutzung in jeder Ableitung eliminieren

läßt. Man kann (II,15, 16) aber auch als Definition einer neuen Aussage $B \rightarrow A$ auffassen, wenn man $\dot{\rightarrow}$ wie bisher als Halbordnung deutet und statt $A \rightarrow B$ jetzt wie in der Verbandstheorie üblich $B \rightarrow A$ schreibt. (II, 16) besagt dann, daß $B \rightarrow A$ die obere Grenze derjenigen Aussagen C ist, die $A \wedge C \dot{\rightarrow} B$ erfüllen (vgl. A II).

Unter Benutzung der expliziten Form (II, 1) der Aussagen folgt, daß

$$(II, 15a) \quad P_A P_{B \rightarrow A} = P_A P_B P_{B \rightarrow A}$$

$$(II, 16a) \quad P_A P_C = P_A P_C P_B \Rightarrow P_C = P_C P_{B \rightarrow A}$$

und da weiterhin

$$(II, 15b) \quad P_A (1 - P_A + P_A P_B) = P_B P_A (1 - P_A + P_A P_B)$$

$$(II, 16b) \quad P_A P_C = P_A P_C P_B \Rightarrow P_C = P_C (1 - P_A + P_A P_B)$$

(so gilt wegen (II, 16a, 16b))

$$(II, 17) \quad P_{B \rightarrow A} = 1 - P_A + P_A P_B$$

Es ist interessant, an Hand von (II,17) den Zusammenhang mit der früher als Halbordnung eingeführten Relation „ $\rightarrow$ “ zu untersuchen. Man sieht nämlich leicht, daß

$$(II, 18) \quad P_A = P_A P_B \Leftrightarrow P_{B \rightarrow A} = 1$$

gilt. $P_{B \rightarrow A}$ ist also der Einheitsoperator, wenn $A \rightarrow B$ gültig ist. Es liegt nahe, diesen Sachverhalt im Rahmen einer „Metalogik“ so zu interpretieren, daß $B \rightarrow A$ dann „wahr“ ist, wenn zwischen A und B die Relation $A \rightarrow B$ besteht. $B \rightarrow A$ ist ja jetzt nicht mehr als eine Relation, sondern als eine Aussage aufzufassen [16].

Verbandstheoretisch bedeutet das, daß $A \rightarrow B$ eine Halbordnungsrelation darstellt, $B \rightarrow A$ dagegen ein Verbandselement, das zu B relative Pseudokomplement von A . Die Existenz eines relativen Pseudokomplements für alle Aussagen hat u. a. die

Distributivität des betreffenden Verbandes zur Folge. Die Beziehung (II, 18) heißt in verbandstheoretischer Schreibweise, daß

$$(II, 19) \quad A \rightarrow B \Leftrightarrow V \rightarrow B \rightarrow A$$

wobei V das Einheitselement des Verbandes ist, das man in relativ-pseudokomplementären Verbänden stets durch $V \Leftrightarrow A \rightarrow A$ definieren kann.

Der Zusammenhang mit der Theorie der linearen Unterräume des Hilbertraumes ist hier der, daß

$$(II, 20) \quad P_{B \rightarrow A} = 1 - P_A + P_A P_B$$

auf den Unterraum $M_B \cup (H - M_A)$ projiziert, also auf alle Teilräume, die aus Linearkombinationen der Elemente von M_B und von $H - M_A$ bestehen. Wir werden hierauf im Zusammenhang mit der Negation noch einmal zurückkommen.

e) Die Negation

Zur Definition der Negation nehmen wir an, daß der Kalkül eine $\wedge$ -Aussage enthält, also eine Aussage, die für alle Aussagen A die Relation $\wedge \rightarrow A$ erfüllt. In den hier untersuchten Kalkülen wird stets eine $\wedge$ -Aussage existieren. Zur Einführung der Negation verwenden wir dann die intuitionistische Definition

$$(II, 21) \quad \neg A \Leftrightarrow A \rightarrow \wedge$$

Auf Grund dieser Definition gelten dann die beiden Sätze:

$$(II, 22) \quad A \wedge \neg A \rightarrow \wedge$$

$$(II, 23) \quad A \wedge C \rightarrow \wedge \Rightarrow C \rightarrow \neg A$$

woraus der Spezialfall

$$(II, 24) \quad A \rightarrow \neg \neg A$$

folgt.

Die explizite Form der Aussage $\neg A$ ist auf Grund der Definition leicht zu finden. Zunächst ist für alle A : $P_{\wedge} = P_{\wedge} P_A$

also $P_{\wedge} = 0$. Nach (II, 21) ist weiter $\neg A = \wedge \neg A$, so daß wegen (II, 20) folgt:

$$(II, 25) \quad P_{\neg A} = 1 - P_A$$

$P_{\neg A}$ projiziert also auf den zu M_A totalsenkrechten Unterraum $H - M_A$. Auf Grund von (II, 25) lassen sich weitere wichtige Sätze über die Negation beweisen, die also über die effektive Logik hinausgehen. Zunächst gilt, daß über (II, 24) hinaus alle Aussagen stabil sind, also

$$(II, 26) \quad \neg \neg A \leftrightarrow A$$

$$\text{da } P_{\neg \neg A} = 1 - (1 - P_A) = P_A.$$

Das tertium non datur

$$(II, 27) \quad V \rightarrow A \vee \neg A$$

ist also für alle Aussagen relativ zulässig. Darüber hinaus kann man (II, 27) hier wegen (II, 14) und (II, 24) auch direkt beweisen, denn es gilt $P_V = 1$ und

$$P_{A \vee \neg A} = P_A + P_{\neg A} - P_A P_{\neg A} = 1$$

Die wichtige verbandstheoretische Bedeutung dieser Beziehung soll im folgenden Abschnitt besprochen werden.

II. 3 Der Logikkalkül

Wir fassen noch einmal die im vorangegangenen Abschnitt gewonnenen Darstellungen der logischen Begriffe zusammen.

A	P_A	M_A
$\neg A$	$1 - P_A$	$H - M_A$
$A \wedge B$	$P_A P_B$	$M_A \cap M_B$
$A \vee B$	$P_A + P_B - P_A P_B$	$M_A \cup M_B$
$A \rightarrow B$	$1 - P_B + P_A P_B$	$M_A \cup (H - M_B)$
$A \rightarrow B$	$P_A \leq P_B$	$M_A \leq M_B$

Für diese Begriffe wurden zunächst im Rahmen der effektiven Logik die folgenden Sätze bewiesen

- (II, 28)
- (1) $A \rightarrow A$
 - (2) $A \rightarrow B, B \rightarrow C \Rightarrow A \rightarrow C$
 - (3) $A \wedge B \rightarrow A$
 - (4) $A \wedge B \rightarrow B$
 - (5) $C \rightarrow A, C \rightarrow B \Rightarrow C \rightarrow A \wedge B$
 - (6) $A \rightarrow A \wedge B$
 - (7) $B \rightarrow A \wedge B$
 - (8) $A \rightarrow C, B \rightarrow C \Rightarrow A \wedge B \rightarrow C$
 - (9) $A \wedge A \rightarrow B \dot{\rightarrow} B$
 - (10) $A \wedge C \rightarrow B \Rightarrow C \dot{\rightarrow} A \rightarrow C$
 - (11) $A \wedge \neg A \rightarrow \wedge$
 - (12) $A \wedge C \rightarrow \wedge \Rightarrow C \rightarrow \neg A$

Unter Benutzung der expliziten Form der Aussagen konnte darüber hinaus gezeigt werden, daß alle Aussagen stabil sind:

$$(II, 29) \quad A \leftrightarrow \neg \neg A$$

und das tertium non datur

$$(II, 30) \quad V \rightarrow A \vee \neg A$$

gilt. Auf Grund der expliziten Form der Aussagen läßt sich auch leicht eine Beziehung zwischen $\rightarrow$ und $\vee$ herleiten. Es ist nämlich

$$(II, 31) \quad B \rightarrow A \leftrightarrow \neg A \vee B$$

da

$$P_{B \rightarrow A} = (1 - P_A + P_A P_B) = P_{\neg A} + P_B - P_{\neg A} P_B = P_{\neg A \vee B}$$

Es soll hier noch etwas auf die verbandstheoretische Struktur des durch (II, 28, 29, 30) charakterisierten Kalküls eingegangen

werden. Bezüglich der Relation $\rightarrow$ bilden die Aussagen eine Halbordnung (II, 28, 1, 2), die zu je zwei Elementen A und B wegen (II, 28, 3, 4, 5) eine untere Grenze $A \wedge B$ besitzt (Halbverband) und wegen (II, 28, 6, 7, 8) eine obere Grenze $A \vee B$, so daß die untersuchte Struktur ein Verband ist. Wie stets in Verbänden gelten daher die zueinander dualen Beziehungen

$$(II, 32) \quad (A \wedge B) \vee (A \wedge C) \rightarrow A \wedge (B \wedge C)$$

$$(II, 33) \quad A \vee (B \wedge C) \rightarrow (A \vee B) \wedge (A \vee C)$$

Die Beziehungen (II, 28, 9, 10) besagen, daß der betrachtete Verband relativ-pseudokomplementär ist, daß also zu je zwei Elementen A und B ein Element $B \rightarrow A$, das zu B relative Pseudokomplement von A , existiert. Es existiert daher zunächst ein Element $V \Leftrightarrow A \rightarrow A$. Weiterhin gilt der wichtige Satz, daß relativ-pseudokomplementäre Verbände distributiv sind, daß also auch die Umkehrungen von (II, 32, 33)

$$(II, 32a) \quad A \wedge (B \vee C) \rightarrow (A \wedge B) \vee (A \wedge C)$$

$$(II, 33a) \quad (A \vee B) \wedge (A \vee C) \rightarrow A \vee (B \wedge C)$$

gelten. Auf den Beweis dieses wichtigen Satzes werden wir in Satz III,1 ausführlich eingehen.

Die für (II, 28, 11, 12) wichtige Annahme, daß eine $\wedge$ -Aussage existiert, bedeutet, daß der betreffende Verband ein Nullelement besitzt. Da der Verband weiter relativ-pseudokomplementär ist, existiert also insbesondere das Pseudokomplement $\rightarrow A$ zu A , das wir oben als Negation bezeichnet haben. Die Beziehungen (II, 28, 11, 12) gelten hier wie stets für Pseudokomplemente.

Darüber hinaus gilt hier aber auch noch (II, 30), weshalb $\rightarrow A$ sogar ein Komplement ist. Da in distributiven Verbänden die Komplementbildung weiterhin eindeutig ist, ist $\rightarrow A$ also das einzige Komplement zu A .

Wir haben daher einen relativ-pseudokomplementären Verband mit 0- und 1-Element, der außerdem (II, 30) erfüllt, also einen komplementären, distributiven Verband, oder auch eine Boolesche Algebra. Die Logik der kommensurablen Eigenschaften ist daher vollkommen äquivalent zur klassischen oder fikti-

ven Logik. Es soll hier noch ganz kurz darauf eingegangen werden, wie sich im Anschluß an die bisherigen Betrachtungen auch die Quantoren $\bigwedge_x$ und $\bigvee_x$ einführen lassen. Wir werden jedoch in den folgenden Kapiteln von den Quantoren keinen weiteren Gebrauch machen.

Wenn in einer Formel eine Aussageform $A(x)$ vorkommt, deren Variable x man nicht spezialisieren darf, da sonst die betreffende Formel ungültig würde, so wollen wir sagen, daß die Variable x in $A(x)$ gebunden ist, und $\bigwedge_x A(x)$ dafür schreiben. x ist dann keine freie Variable mehr und darf nicht mehr spezialisiert werden. Wir definieren dann die Generalisation $\bigwedge_x$ durch

$$(II, 34) \quad \bigwedge_x A(x) \Leftrightarrow \bigwedge_x A(x)$$

Weiterhin führen wir durch die relativ zulässige Regel

$$(II, 35) \quad A(x) \rightarrow \bigvee_x A(x)$$

die Partikularisation $\bigvee_x$ ein. Man kann dann im Rahmen der effektiven Logik zeigen, daß (auf die Beweise wollen wir hier nicht eingehen, vgl. Lorenzen [9])

$$(II, 36) \quad A \wedge \bigvee_x B(x) \leftrightarrow \bigvee_x A \wedge B(x)$$

$$(II, 37) \quad A \vee \bigwedge_x B(x) \leftrightarrow \bigwedge_x A \vee B(x)$$

gilt, und weiterhin

$$(II, 38) \quad \bigvee_x \neg A(x) \rightarrow \neg \bigwedge_x A(x)$$

$$(II, 39) \quad \bigwedge_x \neg A(x) \leftrightarrow \neg \bigvee_x A(x)$$

Mit dem hier gültigen tertium non datur läßt sich darüber hinaus auch noch

$$(II, 40) \quad \neg \bigwedge_x A(x) \rightarrow \bigvee_x \neg A(x)$$

beweisen. Die Beziehungen (II, 36, 37) stellen eine Verallgemeinerung von (II, 32, 33) für unendlich viele Aussagen dar. Verbandstheoretisch entspricht dem die Tatsache, daß die Untermengen des Hilbertraumes zunächst einen vollständigen Verband

bilden, so daß auch für jede unendliche Folge von Teilmengen eine obere und eine untere Grenze existiert. Somit gilt:

$$(II, 41) \quad A \vee \bigwedge_x B(x) \rightarrow \bigwedge_x (A \vee B(x))$$

$$(II, 42) \quad \bigvee_x (A \wedge B(x)) \rightarrow A \wedge \bigvee_x B(x)$$

und die dazu duale Aussage. Da der Verband aber weiterhin relativ pseudokomplementär ist, gelten auch die unendlichen Distributivgesetze (II, 36, 37). Dem entspricht die Tatsache, daß wir uns hier auf die orthogonalen Teilräume beschränkt haben. Für den Verband beliebiger Teilräume gelten, wie im folgenden Kapitel untersucht werden soll, nämlich statt der Distributivgesetze (II, 36, 37) nur noch die stets in vollständigen Verbänden gültigen Beziehungen (II, 41, 42).

III. Die Logik der inkommensurablen Eigenschaften

III. 1 Physik und Logik

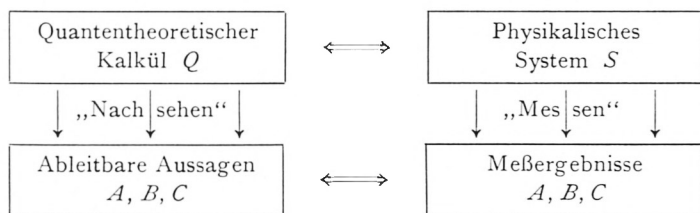
Die im vorangegangenen Kapitel behandelte Logik der kommensurablen Eigenschaften ist eine Logik, die für alle Aussagen gilt, die simultan an einem physikalischen System gemessen werden können. Man könnte daher versucht sein zu meinen, daß damit das Problem der Logik in der Quantentheorie abgeschlossen werden könne, da man sich für alle weiteren Aussagen, die ja nach Voraussetzung nicht mehr gleichzeitig meßbar sind, gar nicht zu interessieren brauchte. Gegen diese Auffassung ist jedoch der Einwand zu erheben, daß es eben eine der bemerkenswerten Eigenschaften der Quantenmechanik ist, daß man an einem gegebenen System jede beliebige Eigenschaft messen kann, wobei man allerdings im allgemeinen die Resultate der vorangegangenen Messung ganz oder teilweise zerstört.

Unter diesen Umständen liegt es nahe, sich nicht auf die Untersuchung einer bestimmten Klasse kommensurabler Eigenschaften zu beschränken, sondern diese zur Klasse aller meßbaren Eigenschaften zu erweitern. Auch wird es oft gar nicht bekannt sein, ob zwei Eigenschaften E_A und E_B kommensurabel sind oder nicht, so daß es sinnvoll ist, von vornherein beliebige Eigenschaften in Betracht zu ziehen. Es soll daher in diesem

Kapitel die Frage untersucht werden, ob sich die für die kommensurablen Eigenschaften gefundenen logischen Regeln ganz oder wenigstens teilweise auf die Klasse beliebiger Eigenschaften übertragen lassen, was dann zu einer „Logik“ allgemein inkommensurabler Eigenschaften führen würde.

Dabei erhebt sich natürlich zunächst die Frage, in welcher Hinsicht ein System solcher Regeln als eine Logik im bisher verwendeten Sinne des Wortes betrachtet werden kann, und mit welchen Methoden sich diese Regeln begründen lassen, bzw. was hier überhaupt unter Begründung zu verstehen ist. Um diese Frage zu diskutieren, ist es nützlich, sich noch einmal zu vergegenwärtigen, wie der Aufbau der Logik der kommensurablen Eigenschaften durchgeführt wurde.

Geht man aus von der Definition $A \Leftrightarrow P_A |f\rangle = |f\rangle$ einer Aussage und läßt nur kommensurable A zu, für die die Projektionsoperatoren vertauschbar sind, so zeigt die Quantentheorie, daß zwischen einer solchen Theorie und der physikalischen Realität sehr einfache Beziehungen bestehen. Der Formalismus Q der Quantenmechanik ist nämlich eineindeutig abbildbar auf die Vorgänge in einem physikalischen System. Die Ergebnisse eines Meßprozesses sind dabei ebenfalls eineindeutig zugeordnet den im Kalkül der Quantenmechanik ableitbaren Aussagen, da hier nur kommensurable Eigenschaften diskutiert werden. Dem „Messen“ einer Eigenschaft E_A an S entspricht daher eineindeutig das „Nachsehen“ in dem Kalkül Q , ob A unter den ableitbaren Aussagen ist.¹ Diese Zusammenhänge können wir durch das folgende Schema verdeutlichen:



¹ Den Kalkül Q denken wir uns aufgebaut aus einigen experimentell gewonnenen Aussagen über den Zustand $|\varphi\rangle$ des untersuchten Systems, die als Anfangsaussagen auftreten, und sämtlichen im Hilbertraum gültigen Beziehungen $A \rightarrow B$, die die Regeln des Kalküls darstellen.

Die Frage, ob für die Meßergebnisse A , B , C die effektive Logik gilt oder nicht, ist durch diese Zuordnung auf das Problem zurückgeführt, ob die effektive Logik für die Aussagen eines Kalküls gilt. Daß dies der Fall ist, kann aber in der operativen Begründung der Logik durch protologische Methoden gezeigt werden. Die Gültigkeit des tertium non datur muß allerdings jeweils an dem speziell vorliegenden Kalkül nachgewiesen werden.

Die gleichen Verhältnisse wie bei kommensurablen Eigenschaften liegen auch in der gesamten klassischen Physik vor, worauf wir jedoch hier nicht näher eingehen wollen. Sowohl in der klassischen Physik als auch in der Physik der kommensurablen Eigenschaften hat daher die Existenz der im obigen Schema angegebenen Abbildung zur Folge, daß für die physikalischen Aussagen die klassische Logik gültig ist. Inwieweit eine Abbildung existiert, muß jeweils von der Physik nachgeprüft werden. Es besteht jedenfalls keinerlei Grund, a priori anzunehmen, daß die Aussagen irgendeines Erfahrungsbereiches eo ipso den Gesetzen der klassischen Logik gehorchen müßten. Wie wir sehen werden, ist dies auch nicht immer der Fall.

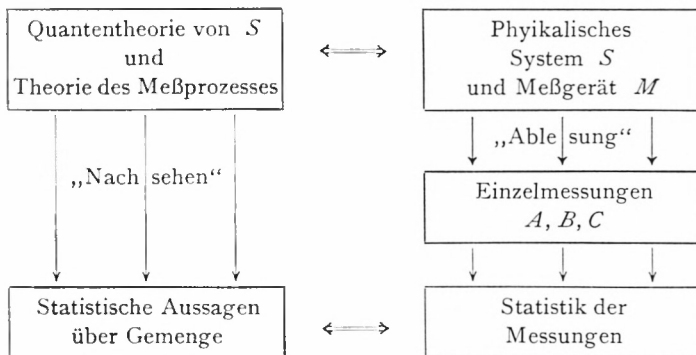
In der Quantenmechanik allgemein inkommensurabler Eigenschaften versagt das obige Schema. Zwar ist es möglich, dem unbeobachteten System eineindeutig den Formalismus der Quantentheorie zuzuordnen, für die Meßergebnisse ist dies jedoch nicht möglich. Die Messung einer beliebigen Größe E_A verändert nämlich den Zustand $|f\rangle$ des Systems in einer solchen Weise, daß der neue Zustand ein Eigenzustand von P_A ist. Da ein entsprechender Vorgang beim „Nachsehen“ in einem Kalkül naturgemäß nicht möglich ist, versagt hier die Zuordnung der Meßergebnisse zu den ableitbaren Aussagen.

Man könnte natürlich daran denken, den Meßapparat M , der diese störenden Wirkungen auf das System ausübt, mit in das System einzubeziehen, so daß $(S + M)$ das neue System ist und die Messung sich auf die reine Ablesung der Resultate reduziert. Dieser letzte Schritt ist insofern unproblematisch, als er kein spezifisch quantentheoretisches Element mehr enthält. Man sieht jedoch leicht, daß man durch ein solches Vorgehen wenig gewinnt (vgl. dazu das untere Schema). Dem System $(S + M)$

hätte man die Theorie von $(S + M)$ zuzuordnen, also die Quantentheorie von S und die Theorie des Meßprozesses. Die Theorie des Meßprozesses macht aber nur die Aussage (vgl. Süßmann [10]), daß der Operator $|f\rangle\langle f|$ des Systems bei der Messung der Observablen A durch die Wechselwirkung des Systems mit dem Meßgerät gemäß der Formel

$$|f\rangle\langle f| \Rightarrow \sum |A_i\rangle\langle A_i|f\rangle\langle f|A_i\rangle\langle A_i|$$

in den statistischen Operator eines Gemenges übergeht. Über den wirklichen Ausgang eines einzelnen Experiments kann die Theorie nichts aussagen. Durch „Nachsehen“ im Kalkül gewinnt man also immer nur statistische Aussagen, die man demgemäß nur mit den statistischen Aussagen über viele Experimente vergleichen kann. Die statistischen Aussagen über viele Experimente sollten daher auch wieder der klassischen Logik gehorchen. Da die Theorie aber über den Ausgang einer einzelnen experimentellen Messung keinerlei Aussagen macht, existiert jetzt auch keine Zuordnung mehr zwischen Meßergebnissen und Kalkülaussagen. Es besteht daher auch kein Grund mehr, anzunehmen, daß beliebige, im allgemeinen inkommensurable Eigenschaften eines Systems S der klassischen Logik gehorchen.



Gehen wir daher von beliebigen, allgemein inkommensurablen Eigenschaften aus. Ebenso wie früher wollen wir Regeln $A \rightarrow B$ untersuchen, die den folgenden Sinn haben: Ist E_A an einem System durch eine Messung nachgewiesen, so hat dieses System

auch die Eigenschaft E_B , d. h. eine Messung wird stets das Vorliegen von E_B ergeben. Unter diesen Regeln und den komplizierteren der Form $C \dot{\rightarrow} A \rightarrow B$ fragen wir wieder nach denjenigen, die unabhängig von den speziellen Aussagen A, B, C stets gelten, d. h. deren Benutzung eliminiert werden kann. Die Frage, wie die Eliminierbarkeit zu beweisen ist, ist jetzt aber neu zu untersuchen. Bei kommensurablen Eigenschaften A, B, C war bekannt, daß diese eineindeutig abbildbar sind auf die ableitbaren Aussagen eines Kalküls. Man konnte daher vergessen, daß man es mit möglichen Meßergebnissen zu tun hat, und sich auf die Diskussion von Kalkülaussagen und deren Logik beschränken. In Kalkülen ist aber auf Grund der protologischen Prinzipien klar, was man unter Eliminierbarkeit zu verstehen hat. Da die eineindeutige Abbildung auf Kalküle jetzt nicht mehr existiert, muß man stets beachten, daß die A, B, C mögliche Meßergebnisse sind. Die Frage, ob gewisse Aussagen gemessen werden können oder nicht, ist bestimmt durch die physikalische Theorie der Aussagen, also die Quantentheorie einschließlich der Theorie des Meßprozesses. Von dieser Theorie brauchen wir hier aber zunächst nur die Behauptung, daß zwei Eigenschaften nur dann simultan gemessen werden können, wenn sie ausdrücklich als kommensurabel vorausgesetzt sind.

Für eine Ableitung, in der allgemeine Aussagen A, B, C und irgendwelche Regeln $A \rightarrow B$ vorkommen mögen und die sich ihrem Wesen nach stets auf ein System in einem bestimmten Zustand bezieht, bedeutet dies: „In einer Ableitung dürfen zwei Aussagen nur dann explizit¹ auftreten, wenn sie ausdrücklich als kommensurabel vorausgesetzt sind.“

Diese Forderung, die wir im folgenden kurz als Komplementaritätsprinzip bezeichnen wollen, besagt also, daß die Inkommensurabilität zweier Aussagen A und B hier ganz im operativen Sinne zu deuten ist: Man kann sie nicht aus einem Experiment gleichzeitig gewinnen und somit auch nicht gemeinsam aufschreiben. Das Komplementaritätsprinzip drückt daher eine starke Einschränkung der operativen Möglichkeiten aus.

¹ Unter „explizitem“ Auftreten einer Aussage A in einer Formel verstehen wir hier, daß eine bestimmte Zeile der Ableitung nur diese Aussage enthält.

Wir können jetzt auch die Frage beantworten, wenn eine Regel $A \rightarrow B$ oder $C \dot{\rightarrow} A \rightarrow B$ als eliminierbar bezeichnet werden kann. Das ist nämlich genau dann der Fall, wenn eine jede Ableitung, die die Regel benutzt, ersetzt werden kann durch eine andere Ableitung, die diese Regel nicht benutzt und die außerdem dem Komplementaritätsprinzip genügt. Die Gesamtheit aller Regeln, Metaregeln usw., die so als allgemeingültig, also für beliebige Aussagen eliminierbar nachgewiesen werden können, wollen wir daher als operative oder auch effektive Quantenlogik bezeichnen.

Vom Gesichtspunkt dieser Logik aus gesehen erscheint das Komplementaritätsprinzip als ein protologisches Prinzip, das hier zu den bisher bekannten protologischen Prinzipien noch hinzugenommen werden muß. Die durch die Physik der Meßaussagen bedingte Einschränkung der operativen Verfügbarkeit über diese Aussagen läßt sich also durch ein neues protologisches Prinzip ausdrücken, welches das Ableiten gewissen Beschränkungen unterwirft.

In der Logik allgemein inkommensurabler Eigenschaften, also der effektiven Quantenlogik, wird man daher die Frage zu untersuchen haben, in welchem Umfang sich die aus der klassischen Logik bekannten Gesetze auch noch unter Beachtung des Komplementaritätsprinzipes beweisen lassen, d. h. ob vielleicht einige dieser Gesetze so allgemein gelten, daß man nicht notwendig die Kommensurabilität der in ihnen enthaltenen Aussagen explizit verlangen muß.

III. 2 Die logischen Begriffe

Der Aufbau der effektiven Quantenlogik, also der effektiven Logik allgemein inkommensurabler Eigenschaften vollzieht sich ähnlich wie der der operativen klassischen Logik. Ausgehend von dem oben besprochenen allgemeinen Aussagetyp, der einem möglichen Experiment entspricht, werden wir untersuchen, welche Regeln, Metaregeln usw. stets eliminierbar sind, wobei aber jetzt zu den bekannten protologischen Prinzipien, mit deren Hilfe die Eliminierbarkeit bisher bewiesen wurde, jetzt noch das Komplementaritätsprinzip hinzugenommen werden muß. Im Ver-

gleich zur Logik kommensurabler Eigenschaften ist also jetzt die Frage zu stellen, welche der in der klassischen Logik gültigen Regeln so allgemein sind, daß sie auch nach Hinzunahme des Komplementaritätsprinzips noch gültig sind, daß man für ihre Gültigkeit also nicht explizit die Kommensurabilität der in ihnen enthaltenen Aussagen fordern muß. Wir werden sehen, daß dies für alle Regeln bis auf eine Ausnahme tatsächlich der Fall ist. Die Gesamtheit dieser Regeln bildet dann die effektive Quantenlogik.

Während das im operativen Aufbau der Quantenlogik benötigte Komplementaritätsprinzip gewissermaßen nur eine begriffliche Minimalforderung darstellt, die man an die Aussagen zu stellen hat, um die effektive Quantenlogik herzuleiten, ist in Wirklichkeit viel mehr über diese Aussagen bekannt, da man ihre explizite Form genau kennt. Aus der Tatsache, daß die Aussagen die Form $A \Leftrightarrow P_A |f\rangle = |f\rangle$ haben, lassen sich daher auch weitere Regeln herleiten, die zwar nicht allgemein, aber speziell für die hier untersuchte Logik quantenmechanischer Aussagen gelten. Das tertium non datur ist von dieser Art. Die um diese speziell für quantenlogische Aussagen gültigen Regeln erweiterte effektive Quantenlogik werden wir im folgenden kurz als Quantenlogik bezeichnen. Diese Quantenlogik ist also nicht allein aus den protologischen Prinzipien zu begründen, da in einige Gesetze spezielle Eigenschaften des verwendeten Aussagetyps wesentlich eingehen.

Im einzelnen werden wir im folgenden Kapitel so vorgehen, daß wir zunächst die logischen Begriffe (III, 2) mit protologischen Mitteln definieren. Gewisse Sätze der effektiven Quantenlogik lassen sich dann sofort beweisen. Die Verwendung der expliziten Gestalt der Aussagen macht es dann weiter möglich, die logischen Begriffe $\wedge$, $\vee$, $\rightarrow$ durch abgeschlossene Linearmannigfaltigkeiten des Hilbertraumes zu interpretieren. Dadurch werden weitere Sätze beweisbar. Dagegen ist die in Kap. II verwendete Darstellung der logischen Partikeln durch zusammengesetzte Projektionsoperatoren hier nicht mehr möglich. Dies liegt im wesentlichen daran, daß sich aus zwei nicht vertauschbaren Projektionsoperatoren keine neuen, nichttrivialen Projektionsoperatoren bilden lassen.

Um die Bedeutung der aus den verschiedenen Regeln jeweils sich ergebenden algebraischen Strukturen besser übersehen zu können, werden wir wiederum die verbandstheoretischen Betrachtungen neben den logischen parallel laufen lassen, ohne im einzelnen jedesmal den Übergang zur anderen Terminologie besonders zu betonen. Die Begründung der logischen Regeln erfolgt natürlich stets mit rein logischen Mitteln. Die verbandstheoretische Terminologie dient nur zur Verdeutlichung der gewonnenen strukturellen Zusammenhänge.

a) Die Halbordnung

Es seien A, B, C beliebige Aussagen, P_A, P_B, P_C die entsprechenden Projektionsoperatoren, so daß etwa A durch

$$(III, 1) \quad A \Leftrightarrow P_A |f\rangle = |f\rangle$$

definiert ist, und M_A, M_B, M_C seien die abgeschlossenen Linear-mannigfaltigkeiten, auf die die Projektionsoperatoren projizieren.

Wir führen dann zwischen zwei Aussagen A und B die Relation $A \rightarrow B$ ein, wobei $A \rightarrow B$ bedeutet, daß wenn eine Eigenschaft A gemessen ist, dann auch B gemessen ist, d. h. daß jede Messung, die eine positive Entscheidung über das Vorliegen von A fällt, auch das Vorliegen von B ergibt. Auf Grund dieser Definition ist klar, daß die Gültigkeit von $A \rightarrow B$ unter anderem besagt, daß A und B kommensurabel sind.

Für die Relation „ $\rightarrow$ “ gelten die beiden Regeln

$$(III, 2) \quad A \rightarrow A$$

$$(III, 3) \quad A \rightarrow B, B \rightarrow C \Rightarrow A \rightarrow C$$

Zum Beweis einer Formel $A \rightarrow B$ hat man zu zeigen, daß im Kalkül der Aussagen A, B, C jede mit Hilfe von $A \rightarrow B$ hergeleitete Aussage auch ohne diese Regel hergeleitet werden kann. Für $A \rightarrow A$ ist das der Fall, wie man sich leicht überzeugt. (III,4) behauptet, daß die Regel $A \rightarrow C$ dann entbehrlich ist,

wenn man die Regeln $A \rightarrow B$ und $B \rightarrow C$ zum Kalkül hinzunimmt. Eine Ableitung, die $A \rightarrow C$ benutzt, sieht folgendermaßen aus

$$\begin{array}{l} 1. \quad A \\ \dots\dots \\ A \rightarrow C \vdash n. \quad C \end{array}$$

Nimmt man die Regeln $A \rightarrow B$ und $B \rightarrow C$ jedoch hinzu, so kann auf $A \rightarrow C$ verzichtet werden, wie die folgende Ableitung zeigt:

$$\begin{array}{l} 1. \quad A \\ \dots\dots\dots \\ A \rightarrow B \vdash n. \quad B \\ B \rightarrow C \vdash n + 1. \quad C \end{array}$$

Benutzt man die explizite Form der Aussagen A, B, C , so bedeutet $A \rightarrow B$, daß immer dann, wenn $P_A |f\rangle = |f\rangle$ gilt, folgt, daß auch $P_B |f\rangle = |f\rangle$ gültig ist, und zwar für jedes $|f\rangle$. Wie in Anh. I gezeigt, ist dafür notwendig und hinreichend, daß

$$(III, 4) \quad P_A = P_A P_B = P_B P_A$$

ist. Das heißt aber, daß P_A und P_B stets vertauschbar sind, also dieselben Verhältnisse vorliegen wie in Kap. II. Der Beweis der Formeln (III, 2, 3) ist unter Benutzung der Operatorschreibweise sehr einfach:

$$(II, 2): \quad P_A^2 = P_A \quad (III, 3): \quad P_A = P_A P_B, \quad P_B = P_B P_C$$

$$\text{also } P_A = P_A P_B P_C = P_A P_C.$$

Bezüglich der abgeschlossenen Linearmanigfaltigkeiten M_A , M_B und M_C ist (III, 4) wiederum äquivalent zu

$$(III, 5) \quad M_A \leq M_B$$

Vom verbandstheoretischen Gesichtspunkt bilden die Aussagen eine Halbordnung bezüglich „ $\rightarrow$ “, wegen (III, 2, 3).

b) Die Konjunktion

Sind A und B Aussagen einer Theorie, so werde eine neue Aussage $A \wedge B$ gemäß der Regel

$$(III, 6) \quad A, B \rightarrow A \wedge B$$

eingeführt. Enthält ein Kalkül wenigstens A und B und nicht das Zeichen $\wedge$, so ist (III, 6) relativ zulässig. Durch Inversion ergeben sich dann wieder die Regeln.

$$(III, 7) \quad A \wedge B \rightarrow A$$

$$(III, 8) \quad A \wedge B \rightarrow B$$

$$(III, 9) \quad C \rightarrow A, C \rightarrow B \Rightarrow C \rightarrow A \wedge B$$

Vom inhaltlichen Gesichtspunkt würde man $A \wedge B$ durch „ A und B “ übersetzen, jedoch sind die folgenden Ableitungen nicht vom inhaltlichen Sinn des Wortes „und“ abhängig. Ohne Bezug auf (III, 6) können wir jetzt (III, 7, 8, 9) als eine Definition der Aussage $A \wedge B$ auffassen, die $A \wedge B$ im Rahmen der hier diskutierten Aussagen eindeutig festgelegt ist. Sei

$$A \wedge B \Leftrightarrow P_{A \wedge B} |f\rangle = |f\rangle$$

die gesuchte Aussage, so ist zunächst wegen (III, 7, 8, 9)

$$(III, 7a) \quad P_{A \wedge B} = P_A P_{A \wedge B}$$

$$(III, 8a) \quad P_{A \wedge B} = P_B P_{A \wedge B}$$

$$(III, 9a) \quad P_A P_C = P_C, P_B P_C = P_C \Rightarrow P_{A \wedge B} P_C = P_C$$

Betrachten wir andererseits den Operator $P_{A \cap B}$, der auf $M_A \cap M_B$ projiziert, so gilt wegen

$$(III, 7b) \quad M_A \cap M_B \leq M_A$$

$$(III, 8b) \quad M_A \cap M_B \leq M_B$$

$$(III, 9b) \quad M_C \leq M_A, M_C \leq M_B \Rightarrow M_C \leq M_A \cap M_B$$

auf Grund der Äquivalenz von (III, 4) und (III, 5) auch

$$(III, 7c) \quad P_{A \cap B} = P_A \cdot P_{A \cap B}$$

$$(III, 8c) \quad P_{A \cap B} = P_B P_{A \cap B}$$

$$(III, 9c) \quad P_C = P_A P_C, \quad P_C = P_B P_C \Rightarrow P_C = P_{A \cap B} P_C$$

wegen (III, 9a) und (III, 7c, 8c) gilt aber

$$P_{A \cap B} = P_{A \wedge B} \cdot P_{A \cap B} = P_{A \cap B} \cdot P_{A \wedge B}$$

und wegen (III, 9c) und (III, 7a, 8a)

$$P_{A \wedge B} = P_{A \cap B} P_{A \wedge B} = P_{A \wedge B} P_{A \cap B}$$

so daß also

$$(III, 10) \quad P_{A \wedge B} = P_{A \cap B}$$

folgt. Damit ist aber $P_{A \wedge B}$ eindeutig bestimmt als der Operator, der auf den Raum $M_A \cap M_B$ projiziert. Eine Darstellung von $P_{A \wedge B}$ als Produkt der Operatoren P_A und P_B ist jedoch jetzt nicht mehr möglich, da $P_A P_B$ im allgemeinen kein Projektionsoperator ist.

In verbandstheoretischer Sprechweise bilden die Aussagen, die (III, 3,4) und (III, 7,8,9) gehorchen, einen Halbverband.

c) Die Disjunktion

Es seien A und B beliebige Aussagen über das System. Durch die Regeln

$$(III, 11) \quad A \rightarrow A \vee B$$

$$(III, 12) \quad B \rightarrow A \vee B$$

führen wir die neue Aussage $A \vee B$ ein. Die Regeln (III, 11, 12) sind in dem Kalkül, der die Aussagen A und B enthält, aber nicht das Zeichen $\vee$, relativ zulässig. Durch Inversion gewinnt man weiter die Regel

$$(III, 13) \quad A \rightarrow C, \quad B \rightarrow C \Rightarrow A \vee B \rightarrow C$$

die in dem Kalkül, der die Aussagen A und B und die Regeln (III, 11, 12) enthält, relativ zulässig ist. Umgangssprachlich würde man $A \vee B$ mit „ A oder B “ übersetzen, ebenso wie in der in Kap. II behandelten Logik. Die im folgenden durchgeführten Ableitungen sind jedoch wiederum nicht von dem inhaltlichen Sinn des Wortes „oder“ abhängig. Es mag deshalb hier genügen, festzustellen, daß das Zeichen $\vee$ ebenso wie in der gewöhnlichen Logik definiert wird und deshalb mit der gleichen Berechtigung wie dort mit „oder“ übersetzt werden kann.

Die Regeln (III, 11, 12, 13) können als eine Definition von $A \vee B$ angesehen werden. Die Aussage $A \vee B$ ist durch (III, 11, 12, 13) eindeutig festgelegt. Sei nämlich

$$(III, 14) \quad A \vee B \Leftrightarrow P_{A \vee B} |f\rangle = |f\rangle$$

so folgt zunächst aus (III, 11, 12, 13)

$$(III, 11 a) \quad P_A = P_{A \vee B} P_A$$

$$(III, 12 a) \quad P_B = P_{A \vee B} P_B$$

$$(III, 13 a) \quad P_A = P_C P_A, \quad P_B = P_C P_B \Rightarrow P_{A \vee B} = P_C P_{A \vee B}$$

Untersucht man andererseits den Operator $P_{A \cup B}$, der auf $M_A \cup M_B$ projiziert, so gilt zunächst auf Grund der Definition von $M_A \cup M_B$

$$(III, 11 b) \quad M_A \leq M_A \cup M_B$$

$$(III, 12 b) \quad M_B \leq M_A \cup M_B$$

$$(III, 13 b) \quad M \leq M_C, \quad M_B \leq M_C \Rightarrow M_{A \vee B} \leq M_C$$

Weiter ist dann wegen der Äquivalenz von (III, 2) und (III, 5)

$$(III, 11 c) \quad P_A = P_{A \cup B} P_A$$

$$(III, 12 c) \quad P_B = P_{A \cup B} P_B$$

$$(III, 13 c) \quad P_A = P_C P_A, \quad P_B = P_C P_B \Rightarrow P_{A \cup B} = P_C P_{A \cup B}$$

Aus (III, 11 a, 12 a) und (III, 13 c) folgt dann

$$P_{A \cup B} = P_{A \vee B} P_{A \cup B} = P_{A \cup B} P_{A \vee B}$$

und aus (III, 11 c, 12 c) und (III, 13 a)

$$P_{A \vee B} = P_{A \cup B} P_{A \vee B} = P_{A \vee B} P_{A \cup B}$$

woraus

$$(III, 14) \quad P_{A \vee B} = P_{A \cup B}$$

folgt. Damit ist der Projektionsoperator $P_{A \vee B}$ eindeutig definiert als der Operator, der auf $M_A \cup M_B$ projiziert.

$$P_{A \vee B} |f\rangle = |f\rangle$$

gilt also genau dann, wenn $|f\rangle \leq M_A \cup M_B$. Eine Darstellung des Operators $P_{A \vee B}$ durch $P_A + P_B - P_A P_B$ wie in Kap. II ist jedoch hier nicht möglich, da $P_A + P_B - P_A P_B$ bei nicht vertauschbaren P_A und P_B kein Projektionsoperator mehr ist.

Die Aussagen, die den Regeln (III, 3, 4), (III, 7, 8, 9) und (III, 11, 12, 13) gehorchen, bilden bezüglich $\wedge$ und $\vee$ einen Verband. Gleichwertig damit ist die Aussage, daß in der Halbordnung (III, 3, 4) für je zwei Elemente A und B eine untere Grenze $A \wedge B$ und eine obere Grenze $A \vee B$ existiert.

d) Die Implikation

In der Logik der kommensurablen Eigenschaften, die sich stets auf die Logik der Kalküle zurückführen läßt, gelten, wie in Kap. II gezeigt wurde, für die Relation „ $\rightarrow$ “ außer den oben eingeführten Halbordnungsrelationen (III, 3, 4) noch die beiden Regeln

$$(III, 15) \quad A \wedge A \rightarrow B \dot{\rightarrow} B$$

$$(III, 16) \quad A \wedge C \rightarrow B \Rightarrow C \dot{\rightarrow} A \rightarrow B$$

Die Regeln (III, 15, 16) konnten dort weiterhin als Definition einer neuen Aussage, der Implikation $B \rightarrow A$ aufgefaßt werden, wenn man $\dot{\rightarrow}$ wie bisher durch die Halbordnung (III, 3, 4) erklärt.

Sind die Eigenschaften A , B , C nicht mehr allgemein kommensurabel, so muß die Gültigkeit von (III, 15, 16) erneut bewiesen werden. Um zu zeigen, daß (III, 15) auch weiterhin gültig ist, muß man zeigen, daß B aus jeder Ableitung eliminiert werden kann, wenn die Aussage A und die Regel $A \rightarrow B$ gegeben ist.

Dies ist aber stets möglich, ohne daß man explizit die Kommen-
surabilität von A und B fordern muß, wie man sich leicht an
einer beliebigen Ableitung, in der B vorkommt, klarmachen
kann.

Verwendet man die explizite Form $A \Leftrightarrow P_A |f\rangle = |f\rangle$ der
quantenlogischen Aussagen, so sieht man sofort, daß aus A und
 $P_A = P_B P_A$ auf $P_B |f\rangle = |f\rangle$ geschlossen werden kann, wie in
(III, 15) behauptet wird.

Die Beziehung (III, 16) ist jedoch nicht mehr allgemeingültig.
Um die Gültigkeit nachzuweisen, hätte man zu zeigen, daß die
Metaregel $C \dot{\rightarrow} A \rightarrow B$ aus jeder Ableitung in einem Metakalkül
eliminiert werden kann, wenn dieser Kalkül die Regel $A \wedge C \rightarrow B$
enthält. Der Beweis dieser Behauptung lautet

1. C
2. $A, C \rightarrow B$
3. $A \rightarrow A$
- 1 $\vdash$ 4. $A \rightarrow C$
- 2, 3, 4 $\vdash$ 5. $A \rightarrow B$ wegen (III, 3) und (III, 9).

Bei diesem Beweis wurde in der vierten Zeile Gebrauch ge-
macht von der Metaregel: $C \dot{\rightarrow} A \rightarrow C$. Zum Beweis hat man zu
zeigen, daß jede Ableitung, die $A \rightarrow C$ benutzt, durch eine an-
dere Ableitung ersetzt werden kann, die $A \rightarrow C$ nicht benutzt,
wenn man C zu den Anfängen des betreffenden Kalküls hinzu-
nimmt. Eine Ableitung, die $A \rightarrow C$ benutzt, hat die Form:

1. A
-
1. $\vdash$ n. C [$A \rightarrow C$]

Die Metaregel $C \dot{\rightarrow} A \rightarrow C$ besagt, daß man auf $A \rightarrow C$ ver-
zichten kann, wenn man C zum Kalkül hinzunimmt, also die
obige Ableitung folgendermaßen umformt:

- o. C ($\vdash C$)
1. A
-
- o $\vdash$ n. C [II, 2]

Man nimmt also C (in Zeile o) hinzu und gewinnt jetzt C in Zeile n durch Verweis auf Zeile o statt auf Zeile 1, unter Benutzung von (III, 2). In dieser Ableitung muß man also annehmen, daß man C wirklich hinzunehmen kann, d. h. daß A und C explizit in derselben Ableitung auftreten können. Nach dem oben formulierten Komplementaritätsprinzip ist das aber nur möglich, wenn A und C kommensurabel sind. Für beliebige Aussagen A und C ist die Hinzunahme von C in der zweiten Ableitung daher nicht möglich. $A \rightarrow C$ ist somit nicht durch eine andere Ableitung zu eliminieren, weshalb $C \dot{\rightarrow} A \rightarrow C$ nicht als gültig nachgewiesen werden kann. Damit wird dann auch der oben angeführte Beweis von $A \wedge C \rightarrow B \Rightarrow C \dot{\rightarrow} A \rightarrow B$ gegenstandslos. Dagegen sind gegen beide Beweise keine Bedenken zu erheben, wenn man A und C als kommensurabel annimmt. Denn über die Kommensurabilität von A mit B und C mit B braucht offensichtlich in den genannten Eliminationsverfahren nichts vorausgesetzt zu werden.

Unter Verwendung der expliziten Form der quantenlogischen Aussagen ist dieser Sachverhalt leicht zu zeigen. Ist $P_A P_C = P_C P_A$, so folgt aus $P_{A \wedge C} = P_B P_{A \wedge C}$ sofort $P_A P_C = P_B (P_A P_C)$ und daraus $P_A P_C = (P_B P_A) P_C$. Das heißt aber, daß für jedes $|f\rangle$, für welches $P_C |f\rangle = |f\rangle$ gilt, auch die Relation $P_A |f\rangle = P_B P_A |f\rangle$, d. h. $A \rightarrow B$ erfüllt ist. Ein entsprechender Beweis für beliebige Aussagen A und C kann jedoch nicht geführt werden.

Die Tatsache, daß (III, 16) nicht mehr allgemein gilt, hat zur Folge, daß man die beiden Beziehungen (III, 15, 16) jetzt auch nicht mehr zu Definition der Implikation $B \rightarrow A$ verwenden kann. Vielmehr wird jetzt eine solche Aussage $B \rightarrow A$, die durch (III, 15, 16) definiert wäre, gar nicht existieren.

Für den Verband der Aussagen $A, B, C \dots$ bedeutet das, daß zu zwei vorgegebenen Aussagen A und B ein relatives Pseudokomplement nur in Ausnahmefällen existiert. Dadurch, daß der Verband der Aussagen jetzt nicht mehr relativ-pseudokomplementär ist, ist natürlich auch die Distributivität des Verbandes in Frage gestellt, die wir in Kap. II aus der Existenz des relativen Pseudokomplementes bewiesen hatten. Die damit zusammenhängenden Fragen sollen im folgenden Abschnitt (III, 3) genauer untersucht werden.

Da im Spezialfall, daß A und C kommensurabel sind, sich (III, 16) als gültig nachweisen läßt, wird es möglich sein, in Kap. III, 3 zu zeigen, daß aus dem eingeschränkten Gesetz (III, 16) auch ein eingeschränktes Distributivgesetz hergeleitet werden kann.

e) Die Negation

Zur Definition der Negation einer Aussage A gehen wir wiederum vom intuitionistischen Begriff der Negation aus und definieren eine Negation $\neg A$ durch

$$(III, 17) \quad \neg A \Leftrightarrow A \rightarrow \wedge.$$

Unter $\wedge$ wollen wir wie bisher eine Aussage verstehen, für die jede Regel $\wedge \rightarrow B$ zulässig ist. In Kap. II sahen wir bereits, daß es stets eine solche $\wedge$ -Aussage gibt, nämlich

$$(III, 18) \quad \wedge \Leftrightarrow P_A |f\rangle = |f\rangle \quad \text{mit} \quad P_{\wedge} = 0.$$

Die in Kap. II untersuchten kommensurablen Eigenschaften bilden einen implikativen Verband, in dem stets ein relatives Pseudokomplement $B \rightarrow A$ zwischen zwei Eigenschaften A und B existiert. Die $\wedge$ -Aussage ist nun mit allen anderen Aussagen kommensurabel, so daß das Pseudokomplement $A \rightarrow \wedge = \wedge \rightarrow A$ stets existiert. Entsprechend (II, 20) ist

$$(III, 19) \quad \neg A \Leftrightarrow P_{\neg A} |f\rangle = |f\rangle$$

mit

$$P_{\neg A} = 1 - P_A + P_A P_{\wedge} = 1 - P_A$$

ebenso wie bei kommensurablen Eigenschaften.

Zunächst gilt, wie stets für ein Pseudokomplement, der Satz vom Widerspruch

$$(III, 20) \quad A \wedge \neg A \rightarrow \wedge$$

und

$$(III, 21) \quad A \wedge C \dot{\rightarrow} \wedge \Rightarrow C \dot{\rightarrow} \neg A$$

Für die speziell hier verwendeten Eigenschaften $P_A |f\rangle = |f\rangle$ gilt ebenso wie in Kap. II, daß die Eigenschaften A wegen der Hermitizität und Idempotenz von P_A stabil sind, also

$$(III, 22) \quad \neg \neg A \leftrightarrow A$$

gilt, weshalb das tertium non datur mit $V \Leftrightarrow A \rightarrow A$

$$(III, 23) \quad V \rightarrow A \vee \neg A$$

in dem Kalkül der Eigenschaften A gilt. Wegen (III, 23) ist daher das Pseudokomplement $\neg A$ sogar ein Komplement. Wegen der Stabilität (III, 22) aller Aussagen gilt weiterhin, daß die Abbildung von A in $\neg A$ einen dualen Automorphismus der Periode 2 (vgl. Anh. II) in dem untersuchten Verband erzeugt, mit den zusätzlichen Eigenschaften (III, 20) und (III, 23). Die Existenz einer solchen Abbildung ist dann wesentlich, wenn in den untersuchten Verbänden kein Distributivgesetz (III, 28, 29) gilt, so daß die Komplementbildung nicht mehr eindeutig ist. Wir werden im folgenden Abschnitt (III, 3) sehen, daß die inkommensurablen Eigenschaften im allgemeinen nur einen modularen Verband bilden. Wegen der Existenz des obenerwähnten Automorphismus ist der Verband jedoch orthokomplementär. Die durch (III, 19) definierte Negation $\neg A$ ist dann das Orthokomplement zu A .

f) Quantoren

Auf die Quantoren wollen wir hier nicht explizit eingehen. Sie lassen sich ebenso wie in Kap. II einführen, gehorchen aber jetzt nur noch den Beziehungen (II, 41, 42), die stets für vollständige Verbände erfüllt sind. Für die folgenden Untersuchungen sind die Quantoren nicht wesentlich.

III. 3 Der Aufbau der Quantenlogik

a) Die effektive Quantenlogik

Im vorigen Abschnitt (III, 2) war gezeigt worden, wie sich die logischen Begriffe für den Fall nicht kommensurabler Eigenschaften definieren lassen. Es hat sich dabei gezeigt, daß, aus-

gehend von diesen Eigenschaften A, B , die den Untermengen M_A und M_B entsprechen, neue Eigenschaften $A \wedge B$, $A \vee B$ und $\neg A$ definieren lassen, – und zwar ebenso wie bei kommensurablen Eigenschaften auf Grund ihrer logischen Bedeutung –, die wiederum eindeutig gewissen Untermengen des Hilbertraumes zugeordnet sind (vgl. dazu Birkhoff und v. Neumann [2]. Das bedeutet, daß die bereits in Kap. II bei kommensurablen Eigenschaften gefundene Möglichkeit, diese Eigenschaften durch Untermengen zu interpretieren, auch hier erhalten bleibt. Dagegen lassen sich die den zusammengesetzten Untermengen entsprechenden Projektionsoperatoren nicht mehr, wie bei kommensurablen Eigenschaften, in einfacher Weise aus den Projektionsoperatoren P_A und P_B zusammensetzen.

Wir stellen die in Kap. III, 2 gefundenen Regeln für die logischen Begriffe $\rightarrow, \wedge, \vee, \neg$ hier noch einmal zusammen.

- (III, 24)
- (1) $A \rightarrow A$
 - (2) $A \rightarrow B, B \rightarrow C \Rightarrow A \rightarrow C$
 - (3) $A \wedge B \rightarrow A$
 - (4) $A \wedge B \rightarrow B$
 - (5) $C \rightarrow A, C \rightarrow B \Rightarrow C \rightarrow A \wedge B$
 - (6) $A \rightarrow A \vee B$
 - (7) $B \rightarrow A \vee B$
 - (8) $A \rightarrow C, B \rightarrow C \Rightarrow A \vee B \rightarrow C$
 - (9) $A \wedge A \rightarrow B \dot{\rightarrow} B$
 - (10) $A \wedge \neg A \rightarrow \wedge$
 - (11) $A \wedge C \rightarrow \wedge \Rightarrow C \rightarrow \neg A$

Die durch die Regeln (III, 24, 1–10) dargestellte Struktur wollen wir als die effektive Quantenlogik bezeichnen, da die Regeln 1–10 allein auf Grund der operativen Bedeutung der logischen Partikeln gelten und hier noch kein Gebrauch von der expliziten Gestalt der verwendeten Aussagen gemacht worden ist. Der wesentliche Unterschied zur Logik der kommensurablen Eigenschaften ist darin zu sehen, daß die Regel

(III, 25) $A \wedge C \rightarrow B \Rightarrow C \dot{\rightarrow} A \rightarrow B$

nicht mehr allgemein gilt, weshalb es nicht mehr möglich ist, hier eine Implikation oder ein relatives Pseudokomplement $B \rightarrow A$ zu definieren.

Diese Einschränkung wird sich beim weiteren Aufbau der Quantenlogik als sehr wesentlich herausstellen, da zahlreiche wichtige Regeln, die in der effektiven Logik beweisbar sind, hier nicht mehr gelten. Die effektive Quantenlogik ist deshalb wesentlich ärmer in der Möglichkeit ihrer Aussagen als die effektive Logik der Kalküle.

Dagegen gelten in der Quantenlogik einige Regeln, die zwar nicht effektiv beweisbar sind, die aber aus der expliziten Gestalt der quantenlogischen Aussagen sich ergeben, d. h. es handelt sich um Regeln, die speziell für den Kalkül der Aussagen

$$A \Rightarrow P_A |f\rangle = |f\rangle$$

zulässig sind. Es wurde bereits bei der Diskussion der Negation darauf hingewiesen, daß die Stabilität $A \leftrightarrow \neg\neg A$ der Aussagen und das tertium non datur $V \rightarrow A \vee \neg A$ von diesem Typ sind. Wir werden im folgenden noch weitere solche Regeln formulieren, die speziell für die quantenmechanischen Aussagen zulässig sind. Daher unterscheidet sich die Quantenlogik von der effektiven Logik auch durch einige Gesetze, die in der Quantenlogik gelten, nicht aber in der effektiven Logik der Kalküle, während in bezug auf (III,25) die umgekehrte Situation vorliegt.

b) Distributivität und Modularität

Auf Grund der Regeln (III, 24, 1–8) bilden die Aussagen der effektiven Quantenlogik einen Verband. Es gelten daher wie stets in Verbänden die Sätze:

$$(III, 26) \quad A \vee (B \wedge C) \rightarrow (A \vee B) \wedge (A \vee C)$$

$$(III, 27) \quad (A \wedge B) \vee (A \wedge C) \rightarrow A \wedge (B \vee C)$$

wobei (III, 27) die zu (III, 26) duale Aussage ist. Gelten nun außerdem noch die Umkehrungen von (III, 26, 27), also

$$(III, 28) \quad (A \vee B) \wedge (A \vee C) \rightarrow A \vee (B \wedge C)$$

$$(III, 29) \quad A \wedge (B \vee C) \rightarrow (A \wedge B) \vee (A \wedge C)$$

so bezeichnet man den Verband als distributiv. Da (III, 28, 29) wieder dual zueinander sind, wird es im folgenden genügen, jeweils nur einen dieser beiden Sätze zu diskutieren. Die Distributivgesetze (III, 28, 29) gelten, wie weiter unten gezeigt wird, in der Quantenlogik nicht allgemein. Es wird aber trotzdem möglich sein, einige interessante Spezialfälle anzugeben, in denen das Distributivgesetz auch noch weiterhin gültig ist.

Die Distributivität eines Verbandes läßt sich aus den beiden zusätzlichen Regeln

$$(III, 30) \quad A \wedge A \rightarrow B \dot{\rightarrow} B$$

$$(III, 31) \quad A \wedge C \rightarrow B \Rightarrow C \dot{\rightarrow} A \rightarrow B$$

beweisen. Da jedoch (III, 31) jetzt nur unter der Voraussetzung gilt, daß A und C kommensurabel sind, so wird auch die Gültigkeit der Distributivgesetze gewissen Einschränkungen unterworfen sein. Um dies zu sehen, wollen wir hier explizit zeigen, wie aus (III, 30, 31) die Distributivität folgt. Jedesmal, wenn (III, 31) in dem Beweis benötigt wird, wird man daher zusätzlich annehmen müssen, daß A und C kommensurabel sind. Durch Zusammenfassung all dieser zusätzlichen Bedingungen gewinnt man dann die Voraussetzungen, unter denen die Distributivität gilt. Wir beweisen also den

Satz (III. 1): Aus (III. 30, 31) folgt

$$A \wedge (B \vee C) \rightarrow (A \wedge B) \vee (A \wedge C)$$

Beweis:

1. $A \wedge (A \rightarrow A \wedge B) \dot{\rightarrow} A \wedge B$
2. $A \wedge B \dot{\rightarrow} (A \wedge B) \vee (A \wedge C)$
- 1, 2 $\vdash$ 3. $A \wedge (A \rightarrow A \wedge B) \dot{\rightarrow} (A \wedge B) \vee (A \wedge C)$
- 3 $\vdash$ 4. $A \rightarrow (A \wedge B) \dot{\rightarrow} A \rightarrow (A \wedge B \dot{\vee} A \wedge C)$ (III, 31)
5. $A \wedge (A \rightarrow A \wedge C) \dot{\rightarrow} A \wedge C$
6. $A \wedge C \dot{\rightarrow} (A \wedge B) \vee (A \wedge C)$
- 5, 6 $\vdash$ 7. $A \wedge (A \rightarrow A \wedge C) \dot{\rightarrow} (A \wedge B) \vee (A \wedge C)$
- 7 $\vdash$ 8. $A \rightarrow A \wedge C \dot{\rightarrow} A \rightarrow (A \wedge B) \vee (A \wedge C)$ (III, 31)

- 4, 8 $\vdash$ 9. $A \rightarrow A \vee B \dot{\vee} A \rightarrow A \vee C \dot{\rightarrow} A \rightarrow (A \wedge B) \vee (A \wedge C)$
 10. $A \rightarrow A \wedge B \dot{\rightarrow} (A \rightarrow A \wedge B) \vee (A \rightarrow A \wedge C)$
 11. $A \rightarrow A \wedge C \dot{\rightarrow} (A \rightarrow A \wedge B) \vee (A \rightarrow A \wedge C)$
 9, 10 $\vdash$ 12. $A \rightarrow A \wedge B \dot{\rightarrow} A \rightarrow (A \wedge B \dot{\vee} A \wedge C)$
 9, 11 $\vdash$ 13. $A \rightarrow A \wedge C \dot{\rightarrow} A \rightarrow (A \wedge B \dot{\vee} A \wedge C)$
 14. $A \wedge B \dot{\rightarrow} A \wedge B$
 14 $\vdash$ 15. $B \dot{\rightarrow} A \rightarrow A \wedge B$ (III, 31)
 14 $\vdash$ 16. $A \dot{\rightarrow} B \rightarrow A \wedge B$
 17. $A \wedge C \dot{\rightarrow} A \wedge C$
 17 $\vdash$ 18. $C \dot{\rightarrow} A \rightarrow A \wedge C$ (III, 31)
 12, 13,
 15, 18 $\vdash$ 19. $B \vee C \dot{\rightarrow} A \rightarrow (A \wedge B \dot{\vee} A \wedge C)$
 20. $A \wedge (B \vee C) \dot{\rightarrow} A \wedge (A \rightarrow A \wedge B \dot{\vee} A \wedge C)$
 20 $\vdash$ 21. $A \wedge (B \vee C) \rightarrow (A \wedge B) \vee (A \wedge C)$

Das Gesetz (III. 31) wurde in den Schritten 4, 8, 15 und 18 benutzt. Wenn also, wie in der effektiven Quantenlogik (III, 31) nur unter der Voraussetzung gilt, daß A und C kommensurabel sind, so muß man also für die Gültigkeit des Distributivgesetzes

$$A \wedge (B \vee C) \rightarrow (A \wedge B) \vee (A \wedge C)$$

die gleichzeitige Entscheidbarkeit von A und $A \rightarrow A \wedge B$; A und $A \rightarrow A \wedge C$; A und B ; A und C voraussetzen. Somit genügt es, die Kommensurabilität von A mit B und A mit C vorauszusetzen. Es ist wesentlich, daß man über die Kommensurabilität von B mit C nichts voraussetzen muß. Wir fassen das Ergebnis noch einmal zusammen:

Satz (III, 2). Ist A mit B und A mit C kommensurabel, so gilt das Distributivgesetz

$$A \wedge (B \vee C) \rightarrow (A \wedge B) \vee (A \wedge C)$$

Während der soeben formulierte Satz sogar in der effektiven Quantenlogik gilt, gilt dies von den im folgenden diskutierten Spezialfällen des Distributivgesetzes nicht mehr. Wir werden

für die Beweise der folgenden Sätze explizit Gebrauch machen von der Tatsache, daß Aussagen stets die Form $P_A |f\rangle = |f\rangle$ haben, wobei P_A ein linearer, hermitischer und idempotenter Operator ist. Für das Distributivgesetz wollen wir jetzt die in (III, 28) angegebene Form verwenden (die ja auf Grund des Dualitätsprinzips äquivalent zu (III, 29) ist), weil sich für diese Form die im folgenden durchzuführenden Beweise etwas durchsichtiger formulieren lassen.

Der hier gemeinte Spezialfall ist der, daß (III, 28) auch in der Quantenlogik gilt, wenn $A \rightarrow B$ oder $A \rightarrow C$ vorausgesetzt werden darf. Das heißt, wir behaupten den folgenden

Satz (III, 3): Ist wenigstens eine der Bedingungen $A \rightarrow B$ und $A \rightarrow C$ erfüllt, so gilt: $(A \vee B) \wedge (A \vee C) \dot{\rightarrow} A \vee (B \wedge C)$.

Beweis: Ist $(A \vee B) \wedge (A \vee C)$, so $|f\rangle \leq M_{A \vee B}$ und $|f\rangle \leq M_{A \vee C}$

also $|f\rangle = |f_A\rangle + |f_B\rangle = |f_{A'}\rangle + |f_C\rangle$

wobei $|f_A\rangle, |f_{A'}\rangle \leq M_A$; $|f_B\rangle \leq M_B$; $|f_C\rangle \leq M_C$

Weiter ist $|f_B\rangle = |f_{A'}\rangle - |f_A\rangle + |f_C\rangle$. Ist $M_A \leq M_C$ so

ist $|f_B\rangle \leq M_C$ also $|f_B\rangle \leq M_B \cup M_C$ und damit

$|f\rangle = |f_A\rangle + |f_B\rangle \leq M_A \cup (M_B \cap M_C)$ wie behauptet.

Zur Durchführung dieses Beweises mußte von der expliziten Gestalt der Aussagen $A \Rightarrow P_A |f\rangle = |f\rangle$ Gebrauch gemacht werden, insbesondere von der Tatsache, daß die Lösungen von $P_{A \vee B} |f\rangle = |f\rangle$ die Form $|f\rangle = |f_A\rangle + |f_B\rangle$ haben, und umgekehrt.

Für die Struktur des Verbandes der Aussagen A, B, C ist der bewiesene Satz insofern von Bedeutung, als er die Modularität des Verbandes garantiert. Auf die wichtigen Konsequenzen dieser Tatsache soll weiter unten eingegangen werden.¹

Der oben im Rahmen der effektiven Quantenlogik bewiesene Satz (III, 2) läßt sich unter Verwendung der in Satz (III, 3) be-

¹ Es sei jedoch bemerkt, daß die Modularität nur für endlich viele Aussagen gilt. Zieht man unendlich viele Aussagen in Betracht, so lassen sich leicht Gegenbeispiele angeben. Vgl. dazu Birkhoff und v. Neumann [2].

nutzten Beweismittel sehr einfach herleiten. Um dies zu zeigen, verwenden wir wie in Satz (III, 3) die Form (III, 28) des Distributivgesetzes, um den Beweis etwas zu vereinfachen.

Satz (III, 4): Ist $P_A P_B = P_B P_A$ und $P_A P_C = P_C P_A$

$$\text{so gilt: } (A \vee B) \wedge (A \vee C) \dot{\rightarrow} A \vee (B \wedge C)$$

Beweis: Ist $|f\rangle \leq M_{A \vee B \wedge A \vee C}$ so gilt: Mit $P_A |f\rangle = |f_A\rangle$
 $|f\rangle = P_{A \vee B} |f\rangle = |f_A\rangle + P_B |f\rangle + P_A P_B |f\rangle =$
 $= |f_A\rangle + |f_B\rangle$
 $|f\rangle = P_{A \vee C} |f\rangle = |f_A\rangle + P_C |f\rangle + P_A P_C |f\rangle =$
 $= |f_A\rangle + |f_C\rangle$
wobei $|f_A\rangle \leq M_A$, $|f_B\rangle \leq M_B$, $|f_C\rangle \leq M_C$ gilt.
Weiter ist $|f_B\rangle = |f_C\rangle$ also $|f_B\rangle = |f_C\rangle \leq M_B \cap M_C$
und damit $|f\rangle = |f_A\rangle + |f_B\rangle \leq M_A \cup (M_B \cap M_C)$
woraus $(A \vee B) \wedge (A \vee C) \rightarrow A \vee (B \wedge C)$ folgt.

Wir haben auf diesen sehr einfachen Beweis in Satz (III, 1) zunächst verzichtet, um zu zeigen, daß Satz (III, 2) sogar in der effektiven Quantenlogik gültig ist.

c) Negation und Komplement

Die Eigenschaften A, B, C bilden, wie oben schon bemerkt, bezüglich der Relationen $\wedge$ und $\vee$ einen Verband, oder auch bezüglich $\rightarrow$ eine Halbordnung, in der für je zwei Elemente A und B stets eine untere Grenze $A \wedge B$ sowie eine obere Grenze $A \vee B$ existiert. Weiterhin sahen wir, daß stets ein Einselement $\vee \rightleftharpoons A \rightarrow A$ sowie ein Nullelement $\wedge$ existiert.

Ein solcher Verband heißt weiter komplementär, wenn es zu jedem A ein A' (das Komplement zu A) gibt, mit den Eigenschaften

$$(III, 32) \quad A \wedge A' = \wedge \quad A \vee A' = \vee$$

Wie in (III, 2) an Hand der Negation gezeigt wurde, gibt es sicher ein Komplement A' , nämlich $A' = \rightarrow A$. Daher ist der

Verband der Aussagen komplementär. Es ist jedoch keineswegs klar, ob $\neg A$ das einzige Komplement zu A ist, oder ob noch weitere Komplemente existieren.

Aus Satz (III. 2) folgt weiter, daß für die Aussagen stets

$$A \rightarrow C \Rightarrow C \wedge (A \vee B) \rightarrow A \vee (B \wedge C)$$

gilt, so daß also der betreffende Verband modular ist. Für die Komplementbildung bedeutet dies, daß der Verband dann auch relativ-komplementär ist. Das heißt, relativ zu zwei vorgegebenen Elementen A, B mit $A \leq B$ gibt es zu jedem Element X ein Element X' , das zu A und B relative Komplement, mit der Eigenschaft

$$(III, 33) \quad X \wedge X' = A \quad X \vee X' = B$$

Weiterhin gilt aber der Satz, daß die relativen Komplemente dann und nur dann eindeutig bestimmt sind, wenn der betreffende Verband distributiv ist. Wie im vorangegangenen Abschnitt gezeigt wurde, kann aber die Distributivität nicht allgemein bewiesen werden, so daß im allgemeinen mehrere relative Komplemente existieren werden. Die hier untersuchten Verbände weisen aber darüber hinaus noch einige Besonderheiten auf, die es dennoch gestatten, in gewissen Fällen die Komplementbildung eindeutig zu machen.

Man bezeichnet einen modularen Verband als orthokomplementär, wenn er einen dualen Automorphismus $X \rightarrow X'$ gestattet, mit den Eigenschaften:

$$(III, 34) \quad (X')' = X \quad X \wedge X' = \wedge \quad X \vee X' = \vee$$

(vgl. Anh. II). Das durch einen solchen Automorphismus definierte Komplement, das Orthokomplement, ist dann eindeutig definiert.

Die Untersuchungen in Kap. III. 2 haben gezeigt, daß der Verband der inkommensurablen Aussagen zwar nicht mehr relativ-pseudokomplementär ist, (woraus die Distributivität folgen würde), daß aber trotzdem zu jedem A ein Pseudokomplement $\neg A$ existiert. Für dieses Pseudokomplement $\neg A$ gelten

aber über die stets für Pseudokomplemente gültige Relation $A \wedge \neg A \rightarrow A$ hinaus, wie bereits bei der Diskussion der Negation gezeigt wurde, die Beziehungen (III, 34), wenn man die explizite Gestalt der quantenlogischen Aussagen mit in Betracht zieht. Das besagt aber, daß die Abbildung $A \rightarrow (\neg A)$ ein Automorphismus im oben diskutierten Sinn darstellt. Das Pseudokomplement ist daher nicht nur ein Komplement, sondern wegen der Stabilität aller Aussagen sogar ein Orthokomplement. Die quantenlogischen Aussagen bilden daher einen orthokomplementären, modularen Verband. Die als Pseudokomplement definierte Negation $\neg A$ einer Aussage A befriedigt (III, 34) und ist daher ein Orthokomplement.

Zu einer beliebigen Aussage A wird es daher im allgemeinen mehrere Aussagen $A, A, \dots$ geben, die den Beziehungen

$$(III, 35) \quad A \wedge A' \leftrightarrow \wedge \quad A \vee A' \leftrightarrow \vee$$

genügen, wenn die Mengen M_A und $M_{A'}$ nicht vertauschbar miteinander sind. Diese Aussagen $A, A', \dots$ sind dann die Komplemente zu A . Sind A und A' jedoch kommensurabel, ist also M_A vertauschbar mit $M_{A'}$, so ist A' die (eindeutig definierte) Negation, für die außer (III, 35) noch $\neg \neg A \leftrightarrow A$ gilt. Die Negation $\neg A$ ist also das einzige zu einer Aussage A kommensurable Komplement. In der Theorie der kommensurablen Eigenschaften sind daher die Verhältnisse wesentlich einfacher, da es dort nur eine Negation gibt, aber keine Komplemente.

Das an der Theorie der inkommensurablen Eigenschaften Neuartige ist daher, daß man aus dem Vorliegen der Beziehung (III, 35) im allgemeinen nicht schließen darf, daß A die Negation zu A sei, wie dies bei kommensurablen Eigenschaften stets gilt.

d) Diskussion einiger Beispiele

Es sollen im folgenden noch einige charakteristische Beziehungen diskutiert werden, die zwar in der klassischen Logik (mit tertium non datur) gelten, nicht aber in der Quantenlogik, und an denen leicht erörtert werden kann, wie stark sich das Fehlen des Distributivgesetzes in der Logik bemerkbar macht.

Betrachten wir zwei inkommensurable Eigenschaften A und B , mit den Unterräumen M_A und M_B . Liegt der Zustand $|f\rangle$ des untersuchten Systems in $|f\rangle \leq M_1 = M_A \cap M_B$, so gilt $A \wedge B$; liegt $|f\rangle$ in dem Unterraum $M_2 = H - (M_A \cap M_B)$, so ist $P_{A \wedge B} |f\rangle = 0$, also gilt dann $\neg(A \wedge B)$. Liegt nun etwa $\neg(A \wedge B)$ vor, so könnte man im Rahmen der klassischen Logik die Regel

$$(III, 36) \quad \neg(A \wedge B) \dot{\rightarrow} \neg A \vee \neg B$$

verwenden, also schließen, daß dann auch $\neg A$ oder $\neg B$ gelten muß. (III, 36) ist aber in der Quantenlogik nicht mehr gültig. (Die Umkehrung gilt auch weiterhin). Der Grund ist natürlich der, daß M_A und M_B nicht vertauschbar sind, so daß $|f\rangle \leq M_2$ liegen kann, ohne deswegen in $(H - M_A) \cup (H - M_B)$ liegen zu müssen wie bei vertauschbaren Unterräumen. Setzt man außer $\neg(A \wedge B)$ noch voraus, daß $P_A |f\rangle = |f\rangle$ gilt, so würde man in der klassischen Logik das dort gültige Gesetz

$$(III, 37) \quad A \wedge \neg(A \wedge B) \dot{\rightarrow} \neg B$$

anwenden und somit auf $\neg B$ schließen können. (III, 37) gilt aber wiederum nicht in der Quantenlogik, so daß man hier nicht auf $\neg B$ schließen kann. Daraus folgt jedoch keineswegs, daß B vorliegt. Vielmehr wird man weder B noch $\neg B$ beweisen können.

Abgesehen von diesen eben erörterten Fällen wollen wir jetzt annehmen, daß $|f\rangle$ nicht in einer der oben diskutierten Mengen M_1, M_2 , liegt. Es wird dann weder $P_{A \wedge B} |f\rangle = |f\rangle$ noch $P_{A \wedge B} |f\rangle = 0$ sein, so daß also weder $A \wedge B$ noch $\neg(A \wedge B)$ gültig ist. Nehmen wir nun weiter an, daß $P_A |f\rangle = |f\rangle$ für den untersuchten Zustand gilt. In der klassischen Logik könnte man dann das Gesetz

$$(III, 38) \quad A \dot{\rightarrow} (A \wedge B) \vee (A \wedge \neg B)$$

verwenden, welches aber in der Quantenlogik nicht gilt. (Die Umkehrung gilt weiterhin.) (III, 38) entspricht etwa dem, was oft etwas unpräzise als *tertium non datur* bezeichnet wird: Daß nämlich eine Eigenschaft B vorliegt oder nicht, wenn man be-

reits weiß, daß A gilt.¹ Man sollte hier vielleicht präziser von einem tertium non datur relativ zu A sprechen. Während das relative tertium non datur (III,38) unter den genannten Voraussetzungen nicht gültig ist, ist jedoch das absolute tertium non datur $B \vee \neg B$ wie oben gezeigt wurde, auch in der Quantenlogik stets für alle $|f\rangle$ erfüllt.

IV. Der quantenlogische Modalkalkül

IV. 1 Sprache und Metasprache

Die Logik beliebiger, im allgemeinen nicht kommensurabler Eigenschaften war in Kap. III als eine Objektsprache entwickelt worden. Es hat sich dabei herausgestellt, daß man auf Grund der allgemein gültigen Regeln der Quantenlogik für ein vorgegebenes System $|f\rangle$ und eine beliebige Eigenschaft A im allgemeinen weder A noch $\neg A$ beweisen kann. Da andererseits die Eigenschaften A und $\neg A$ sehr wohl in einem Meßprozeß als vorliegend sich erweisen können (allerdings dann nicht in dem Zustand $|f\rangle$), liegt es nahe, diesen Sachverhalt mit Hilfe der Modalitäten zu interpretieren: Die Eigenschaften A und $\neg A$ sind „möglich“ in $|f\rangle$. „Notwendig“ würde man eine Eigenschaft A dagegen nennen, wenn man mit Sicherheit sagen kann, daß ein Meßprozeß das Vorliegen von A ergibt.

Eine solche Ausdrucksweise würde am ehesten dem Sprachgebrauch innerhalb der Physik entsprechen (vgl. dazu auch v. Weizsäcker [11], [12]). Das liegt daran, daß dort die Meßergebnisse als Objekte behandelt werden und jede physikalische Interpretation ein Sprechen über diese Objekte ist, d. h. es wird gar nicht der Kalkül der Objekte betrachtet, sondern es wird ihre Meßbarkeit, ihre Möglichkeit und ihre Unmöglichkeit diskutiert. In der Sprache der Logik heißt das, daß nicht die Objekt-

¹ Man kann (III, 38) noch dahingehend verschärfen, daß $P_A = |f\rangle \langle f|$ nur auf den untersuchten Zustand projiziert. Die Ungültigkeit von (III, 38) für die hier untersuchten $|f\rangle$ besagt dann, daß für den vorgegebenen Zustand $|f\rangle$ weder $P_B|f\rangle = |f\rangle$ noch $P_{\neg B}|f\rangle = |f\rangle$ gilt, B also in bezug auf $|f\rangle$ unbestimmt ist.

logik und die in ihr gültigen Regeln diskutiert werden, sondern daß in einer Metasprache das Ableiten, d. h. hier also die Meßbarkeit, selbst diskutiert wird. Wir werden im folgenden zeigen, daß eine solche Metasprache sich sehr durchsichtig formulieren läßt, wenn man von den obenerwähnten Modalitäten in geeigneter Weise Gebrauch macht.

Bevor wir mit der Formulierung einer solchen Metasprache beginnen, ist es aber nützlich, sich zu vergegenwärtigen, wie sich die Modalitäten in der operativen Logik der Kalküle einführen lassen [9]. Wir gehen dazu von einem Kalkül K aus, der durch die Zeichen $\rightarrow, \wedge, \vee, \neg$ erweitert werden möge. Die für die Fragestellung der Modallogik besondere Situation liegt nur dann vor, wenn man über den Kalkül K nicht alles weiß, sondern wenn nur eine endliche Klasse von Formeln aus K bekannt ist. Die Konjunktion W all dieser bekannten Formeln möge das „Wissen“ heißen. Wir wollen nun eine Aussage A notwendig relativ zu dem Wissen W nennen, wenn A aus W ableitbar ist, und möglich (relativ zu W), wenn $\neg A$ nicht aus W ableitbar ist. Wir definieren also

$$(IV, 1) \quad \Delta_W A \Leftrightarrow W \vdash A$$

$$(IV, 2) \quad \nabla_W A \Leftrightarrow W \not\vdash \neg A$$

Der Kalkül der Metaregeln für Δ_W und ∇_W ist damit zurückgeführt auf den Kalkül der Metaaussagen $A \vdash B$. In der Modallogik untersucht man nun diejenigen Regeln dieses Metakalküls, die unabhängig von der speziellen Wahl von K und W gültig sind. Die Ableitbarkeit irgendwelcher Metaaussagen $A \vdash B$ ist aber durch die effektive Logik der Kalküle definiert. Das bedeutet, daß man die Regeln des Modalkalküls gewinnt auf Grund der allgemeinen Einsicht, die man über das Ableiten von Aussagen in Kalkülen in der operativen Logik gewonnen hat. Denn unabhängig davon, ob man K kennt oder nicht, und wie K im einzelnen beschaffen ist, weiß man doch, daß K ein Kalkül ist, und somit die in der effektiven Logik gültigen Gesetze über die Ableitbarkeit von Aussagen auch für K gültig sind.

Eine in vieler Hinsicht ähnliche Situation wie in der Modallogik liegt nun auch in der Metalogik quantenmechanischer Aus-

sagen vor. Eine bestimmte Eigenschaft eines physikalischen Systems S wird allerdings nicht in einem Kalkül abgeleitet, sondern an dem betreffenden System gemessen. Das uns jeweils maximal zur Verfügung stehende Wissen W über das System S ist der Zustand $|f\rangle$, in dem sich S gerade befindet. Da es weiterhin nicht möglich ist, aus der Kenntnis von $|f\rangle$ den Meßwert einer beliebigen Observablen vorauszusagen, das Wissen W also niemals vollständig ist, liegt es nahe, auch hier von den Modalitäten Gebrauch zu machen: Eine Eigenschaft A möge notwendig relativ zu W heißen, wenn A mit Sicherheit bei einer Messung an $|f\rangle$ als zutreffend sich erweist, und möglich, wenn $\neg A$ nicht mit Sicherheit für einen Meßprozeß vorausgesagt werden kann. Statt dessen wollen wir auch sagen, daß A aus W hergeleitet werden kann, bzw. $\neg A$ nicht aus W abgeleitet werden kann, um uns an den in der Quantenlogik (Kap. III) eingeführten Sprachgebrauch anzulehnen. Wir definieren also:

$$(IV, 3) \quad \Delta_W A \Leftrightarrow W \vdash A$$

$$(IV, 4) \quad \nabla_W A \Leftrightarrow W \not\vdash \neg A$$

Die Modalitäten sind damit wiederum auf Metaaussagen $A \vdash B$ zurückgeführt. Als „quantenlogischen Modalkalkül“ wollen wir jetzt den Metakalkül all der Regeln bezeichnen, die unabhängig von der speziellen Wahl von W sind.

Die Regeln über die Ableitbarkeit von Aussagen konnten in der klassischen Modallogik gewonnen werden auf Grund der Kenntnisse über das Ableiten in Kalkülen. Hier dagegen ist kein Kalkül gegeben, sondern ein quantenmechanisches System, bzw. ein Wissen W über dieses System. Die Ableitbarkeit der Metaaussagen $A \vdash B$ ist daher nicht mehr durch die effektive Logik der Kalküle definiert, sondern durch die in Kap. III diskutierte effektive Quantenlogik bzw. die hier immer gültige Erweiterung zur Quantenlogik. Die Quantenlogik spielt hier also für den quantenlogischen Modalkalkül dieselbe Rolle wie die effektive Logik der Kalküle für die klassische Modallogik.

Damit ist die Methode für den Aufbau eines quantenlogischen Modalkalküls klar. Ausgehend von den Definitionen (IV, 3, 4) und den aus der effektiven Quantenlogik folgenden Sätzen über

die Metaaussagen $A \vdash B$ werden wir einen Kalkül für Δ_W und ∇_W aufstellen. Durch diesen Kalkül werden dann die Zeichen Δ_W und ∇_W charakterisiert sein, so daß wir nachträglich wieder auf die Definition (IV, 3, 4) verzichten können.

Im Anschluß an den Aufbau dieses Kalküls sollen einige besondere Eigenschaften dieses Formalismus genauer diskutiert werden. Während der klassische Modalkalkül nur dann eine Bedeutung hat, wenn man zufällig den untersuchten Kalkül nicht genügend kennt, ist der quantenlogische Modalkalkül von viel grundsätzlicherer Bedeutung. Die Tatsache, daß man niemals mehr über ein physikalisches System wissen kann als seinen Zustand $|f\rangle$, und daß dieses Wissen immer unvollständig ist, hat zur Folge, daß man auch hypothetisch nicht auf die Modalitäten verzichten kann. Wir werden diese Probleme an Hand der Begriffe der Wirklichkeit, Objektivität und Koexistenz weiter unten genauer diskutieren. Hierbei werden die wesentlichen Unterschiede zwischen dem klassischen und dem quantenlogischen Modalkalkül sichtbar werden.

IV. 2 Der Aufbau des Modalkalküls

Wir gehen aus von einem physikalischen System S und dem Kalkül der Aussagen A, B, C über dieses System, der durch die logischen Partikeln $\rightarrow, \wedge, \vee, \neg$ erweitert sein möge. Über das System S besitzen wir ein Wissen W , was maximal durch den Zustand $|f\rangle$ des Systems bzw. durch $P_W = |f\rangle\langle f|$ gegeben ist. Im allgemeinen ist W die Konjunktion aller über S bekannten Eigenschaften A_i . Relativ zu W definieren wir dann die Modi Δ_W (A ist notwendig relativ zu W) und ∇_W (A ist möglich relativ zu W) durch (IV, 3, 4). Die Ableitbarkeit der Metaaussagen $A \vdash B$ ist durch die in Kap. III diskutierte effektive Quantenlogik (III, 24) definiert. Benutzen wir $\Rightarrow, \bar{\wedge}, \bar{\vee}, \Rightarrow$ als logische Partikeln im Modalkalkül und fügen wir noch $\underline{\vee}$ (das Wahre) und $\bar{\wedge}$ (das Falsche) hinzu, so gelten die Regeln:

- (IV, 5)
- | | |
|---|-----------------------------|
| 1 | $A \vdash \underline{\vee}$ |
| 2 | $\bar{\wedge} \vdash B$ |
| 3 | $A \vdash A$ |

- 4 $A_1 \vdash A_2 \bar{\wedge} A_2 \vdash A_3 \Rightarrow A_1 \vdash A_3$
- 5 $A \bar{\wedge} A \rightarrow B \vdash B$
- 6 $A \vdash B_1 \bar{\wedge} A \vdash B_2 \Leftrightarrow A \vdash B_1 \wedge B_2$
- 7 $A_1 \vdash B \bar{\wedge} A_2 \vdash B \Leftrightarrow A_1 \vee A_2 \vdash B$
- 8 $A \vdash \neg B \Rightarrow A \wedge B \vdash \wedge$
- 9 $\bar{\wedge}_x A \vdash B(x) \Leftrightarrow A \vdash \bigwedge_x B(x)$
- 10 $\bar{\wedge}_x A(x) \vdash B \Leftrightarrow \bigwedge_x A(x) \vdash B$

Bei (9) komme x nicht frei in A vor und bei (10) nicht frei in B .

Der Unterschied dieses quantenlogischen Metakalküls zum klassischen Metakalkül besteht darin, daß die Formel

$$(IV, 6) \quad A \wedge B \vdash C \Rightarrow A \vdash B \rightarrow C$$

nur gilt, wenn A und B kommensurabel sind, und demgemäß auch die Umkehrung von (IV, 5(8)) nur in diesem Spezialfall gültig ist. Die Negation einer Metaaussage $\Rightarrow \mathfrak{A}$ definieren wir wieder durch $\mathfrak{A} \Rightarrow \bar{\wedge}$. Dann ist $A \not\vdash B$ im (IV, 4) als Abkürzung von $\Rightarrow A \vdash B$ aufzufassen.

Die Erweiterung der effektiven Quantenlogik zur Quantenlogik erfolgt hier durch Hinzunahme des tertium non datur sowohl in der Objektsprache

$$(IV, 7) \quad \vee \vdash A \vee \neg A$$

als auch in der Metasprache

$$(IV, 8) \quad \mathfrak{A} \vee \Rightarrow \mathfrak{A}$$

Die Gültigkeit von (IV, 7) in der Quantenlogik war im Kap. III eingehend diskutiert worden. Ebenso läßt sich (IV, 8) unter Verwendung des Kalküls der Projektionsoperatoren beweisen. Eine Metaaussage $\mathfrak{A} \Leftrightarrow A \vdash B$ heißt dann $P_A P_B = P_B$, und $\Rightarrow \mathfrak{A}$ demgemäß $P_A P_B \neq P_B$. Da für Gleichheitsrelationen aber das tertium non datur stets zulässig ist, folgt damit die Zulässigkeit von (IV, 8) (vgl. Lorenzen [9]).

Betrachten wir jedoch zunächst den effektiven quantenlogischen Modalkalkül, der sich ohne Benutzung von (IV, 7, 8) ergibt. Es gelten dann die Regeln:

- (IV, 9) 1 $\Rightarrow \Delta_W \wedge$
 2 $A \vdash B \bar{\wedge} \Delta_W A \Rightarrow \Delta_W B$
 3 $\Delta_W A \bar{\wedge} \Delta_W B \Rightarrow \Delta_W (A \wedge B)$
 4 $\Delta_W \vee$
 5 $\nabla_W A \Leftrightarrow \Delta_W \rightarrow A$
 6 $\bar{\bigwedge}_x \Delta_W A(x) \Rightarrow \Delta_W \bar{\bigwedge}_x A(x)$

(1) besagt, daß W keine $\wedge$ -Aussage ist. (2) und (3) folgen direkt aus (IV, 3, 4). (5) zeigt, daß ∇_W durch Δ_W definierbar ist.

Der durch (IV, 5) und (IV, 9) gebildete Kalkül möge als effektiver quantenlogischer Modalkalkül bezeichnet werden. Die Modi Δ_W und ∇_W sind jetzt durch (IV, 5, 9) charakterisiert. Ist nämlich W eine endliche oder unendliche Konjunktion von Aussagen A , so gilt wegen (IV, 9: 2, 3, 6)

$$(IV, 10) \quad W \vdash B \bar{\wedge} \bar{\bigwedge}_x \Delta A(x) \Rightarrow \Delta B$$

Betrachtet man nun in einem etwas verallgemeinerten Sinne die unendliche Konjunktion W aller Aussagen A , für welche ΔA ableitbar ist, so gilt wegen (IV, 10)

$$(IV, 11) \quad \Delta B \Leftrightarrow W \vdash B$$

und weiter wegen (IV, 9(5))

$$(IV, 12) \quad \nabla B \Leftrightarrow W \not\vdash \rightarrow B$$

wodurch jetzt die Definitionen (IV, 3, 4) überflüssig werden.

Wir wollen hier nicht im einzelnen darauf eingehen, welche Sätze im quantenlogischen Modalkalkül hergeleitet werden können. Diejenigen Eigenschaften und Sätze des Kalküls, die für die Interpretation der Quantenmechanik wesentlich sind, sollen jedoch im folgenden Abschnitt (Kap. IV, 3) besprochen werden.

Wir wollen hier nur noch eingehen auf einige Beziehungen zwischen den Operatoren Δ_W und ∇_W und den durch Anwendung

der Negation daraus entstehenden Operatoren. Die zunächst möglichen 9 Operatoren aus Δ :

$$\Delta, \Delta \neg, \Delta \neg \neg; \Rightarrow \Delta, \Rightarrow \Delta \neg, \Rightarrow \Delta \neg \neg; \Rightarrow \Rightarrow \Delta, \Rightarrow \Rightarrow \Delta \neg, \\ \Rightarrow \Rightarrow \Delta \neg \neg$$

und die 4 Operatoren aus ∇ :

$$\nabla, \nabla \neg; \Rightarrow \nabla, \Rightarrow \nabla \neg$$

(für ∇ gibt es keine weiteren, da wegen (IV, 9(5)) $\nabla \neg \neg = \nabla$ und $\Rightarrow \Rightarrow \nabla = \nabla$ gilt) lassen sich durch Anwendung des tertium non datur (IV, 7, 8), welches hier stets gilt, auf die 4 Modi

$$(IV, 13) \quad \begin{array}{ll} \Delta = \Rightarrow \nabla \neg & \Rightarrow \Delta = \nabla \neg \\ \nabla = \Rightarrow \Delta \neg & \Rightarrow \nabla = \Delta \neg \end{array}$$

reduzieren. Dadurch ist jetzt Δ durch ∇ definierbar geworden. Zwischen diesen 4 Modi gelten weiterhin die Relationen

$$(IV, 14) \quad \Delta A \Rightarrow \nabla A$$

wegen $\Delta A \bar{\wedge} \Delta \neg A \Rightarrow \Delta(A \wedge \neg A) \Rightarrow \Delta \wedge \Rightarrow \bar{\wedge}$
nach (IV, 9, 1, 3) und

$$(IV, 15) \quad \Rightarrow \nabla \Rightarrow \Rightarrow \Delta$$

IV. 3 Möglichkeit und Wirklichkeit

In der klassischen Modallogik (Becker [17], Lewis-Lanford [18], Lorenzen [9]) ist es üblich, noch den Fall zu betrachten, in dem die Modi sich sogar auf zwei reduzieren lassen. Wir wollen hierauf noch kurz eingehen, da sich an Hand dieser Theorie die wesentlichen Unterschiede zwischen dem klassischen und dem quantenlogischen Modalkalkül besonders deutlich zeigen werden.

In der klassischen Modallogik untersucht man dazu ein Wissen W_0 (W_0 möge wieder die Konjunktion mehrerer Aussagen sein) mit $W_0 \not\vdash \wedge$ und

$$(IV, 16) \quad W_0 \vdash A \vee W_0 \vdash \neg A$$

für alle Aussagen A . Ein solches Wissen, aus dem somit für jede Aussage A entweder A oder $\neg A$ hergeleitet werden kann, wollen wir ein vollkommenes Wissen nennen. Häufig wird auch W_0 als „Wirklichkeit“ bezeichnet im Gegensatz zu dem nur bruchstückhaften Wissen W , das wir von dieser Wirklichkeit haben. Bezüglich W_0 fallen die Modi Δ_{W_0} und ∇_{W_0} zusammen, d. h. es gilt

$$(IV, 17) \quad \Delta_{W_0} A \Leftrightarrow \nabla_{W_0} A$$

Es ist daher angebracht, eine neue Bezeichnung einzuführen. Wir definieren

$$\Diamond_{W_0} A \equiv W_0 \vdash A$$

und sagen, A sei „wirklich“ relativ zu W_0 .

Auf Grund der angegebenen Bedeutung der Wirklichkeit W_0 ist klar, daß ein beliebiges unvollständiges Wissen W mit W_0 in der Beziehung $W_0 \vdash W$ stehen muß. Daraus ergibt sich eine Beziehung zwischen den drei Modalitäten, das sogen. Modalgefälle:

$$(IV, 18) \quad \Delta_W A \Rightarrow \Diamond_{W_0} A \Rightarrow \nabla_W A$$

Eine völlig andere Situation liegt nun in dem quantenlogischen Modalkalkül vor. Das Maximum an Wissen, das man über ein vorgelegtes System S haben kann, ist immer der Zustand $|f\rangle$ dieses Systems. Ein solches Wissen W hat aber nie die Eigenschaft, daß für alle Aussagen A

$$(IV, 19) \quad W \vdash A \vee W \vdash \neg A$$

gilt, sondern dies ist nur für diejenigen A richtig, die mit W kommensurabel sind. Eine Eigenschaft A , die (IV, 19) gehorcht, wollen wir hier „objektiv“ in bezug auf W nennen. Für objektive Aussagen gilt dann wieder, daß die relativen Modi Δ_W und ∇_W zusammenfallen, also:

$$(IV, 20) \quad \Delta_W A \Leftrightarrow \nabla_W A.$$

Dagegen gilt hier eine etwas schwächere Aussage als (IV, 19) auch für alle Aussagen. Ist W ein maximales Wissen, also

$P_W = |f\rangle \langle f|$, so gilt $[P_W, P_A] = 0 \vee [P_W, P_A] \neq 0$, und da weiterhin $[P_W, P_A] = 0$ äquivalent ist zu $\Delta_W A \vee \Delta_W \neg A$, so folgt wegen (IV, 13)

$$(IV, 21) \quad \Delta_W A \dot{\vee} \Delta_W \neg A \dot{\vee} (\nabla_W A \wedge \nabla_W \neg A).$$

Es liegt nahe, (IV, 19) als ein tertium non datur relativ zu W zu bezeichnen, da (IV, 19) aussagt, daß bei Vorliegen von W entweder A oder $\neg A$ ableitbar ist. (IV, 19) ist die metasprachliche Formulierung des bereits in (III, 38) in der Objektsprache ausgedrückten Sachverhalts, daß bei Kenntnis von W für jede mit W kommensurable Eigenschaft entweder $W \wedge A$ oder $W \wedge \neg A$ beweisbar ist. Davon streng zu unterscheiden sind die immer, auch für beliebige Aussagen, gültigen Gesetze (IV, 7, 8), die üblicherweise als tertium non datur in der Objekt- bzw. Metasprache bezeichnet werden.

Die ebenfalls für beliebige A gültige Formel (IV, 21) könnte man in dieser Redeweise als eine metasprachliche Formulierung des „quantum non datur“ relativ zu W bezeichnen. Denn (IV, 21) drückt aus, daß es für eine beliebige Eigenschaft A in der Quantenlogik höchstens drei Möglichkeiten gibt, daß nämlich $\Delta_W A$, $\Delta_W \neg A$ oder $\nabla_W A \wedge \nabla_W \neg A$ erfüllt ist. In diesem Sinne kann man sagen, daß in der Quantenlogik das (relative) tertium non datur (IV, 19) zwar nicht mehr allgemein gilt, daß dafür aber das (relative) quantum non datur (IV, 21) stets erfüllt ist. Dagegen gilt hier natürlich weiterhin für alle Aussagen der (relative) Satz vom Widerspruch. (Wegen (IV, 9: 1, 3))

$$(IV, 22) \quad \Delta_W A \bar{\wedge} \Delta_W \neg A \Rightarrow \bar{\wedge}.$$

Ist für eine vorgegebene Aussage A nicht gerade $\Delta_W A$ oder $\Delta_W \neg A$ beweisbar, sind also A und W inkommensurabel, so folgt aus (IV, 21), daß dann

$$(IV, 23) \quad \nabla_W A \bar{\wedge} \nabla_W \neg A$$

gilt. Dies ist der für die quantenlogische Situation am meisten charakteristische Fall. (IV, 23) besagt nämlich, daß zwei im klassischen Sinne kontradiktorische Aussagen A und $\neg A$ hier bei-

de gleichzeitig möglich sind, wenn ein zu A inkommensurables Wissen W vorliegt. Wir wollen dann sagen, daß die kontradiktorischen Eigenschaften A und $\neg A$ sich im Modus der „Koexistenz“ relativ zu W befinden (vgl. dazu v. Weizsäcker [12]). Die Koexistenz ist also ebenso wie die Modi der Möglichkeit und der Notwendigkeit ein relativer Modus, so daß man schärfer vom Begriff der „relativen Koexistenz“ sprechen sollte. Es ist jedoch auf Grund der Theorie des quantenmechanischen Meßprozesses klar, daß jedes Experiment am System S , das eine Entscheidung über die Eigenschaft A zum Ziele hat, die Koexistenz zerstören muß: Es würde sich dann entweder A oder $\neg A$ ergeben, da durch den Meßprozeß der Zustand des Systems und damit auch das Wissen W in einer solchen Weise verändert wird, daß das neue Wissen mit A kommensurabel ist.

Sind A und W jedoch kommensurabel, so ist A eine „objektive“ Eigenschaft des Systems im obigen Sinne. Es gilt dann für A die Beziehung (IV, 19), woraus weiter

$$(IV, 24) \quad \nabla_W A \bar{\wedge} \nabla_W \neg A \Rightarrow \bar{\wedge}$$

folgt. (IV, 24) besagt aber, daß die Koexistenz einer objektiven Eigenschaft A mit ihrem kontradiktorischen Gegenteil $\neg A$ unmöglich ist.

An Hand des in (IV, 16) eingeführten Begriffes der Wirklichkeit kann man sich noch leicht klarmachen, wie entscheidend sich die Hinzunahme der in dem quantenlogischen Modalkalkül nicht mehr gültigen Beziehung (IV, 6) bemerkbar machen würde. Ist W ein Wissen und W_0 die Wirklichkeit, so gilt in der klassischen Modallogik, in der (IV, 6) gültig ist, der Satz:

$$(IV, 25) \quad \nabla_W A \Leftrightarrow \bigvee_{W_0} \Diamond_{W_0} W \bar{\wedge} \Diamond_{W_0} A$$

(auf den Beweis wollen wir hier nicht eingehen, vgl. dazu Lorenzen [9]). (IV, 25) besagt, daß, wenn eine Aussage A bezüglich W möglich ist, es dann stets ein vollständiges Wissen W_0 gibt, derart, daß sowohl W relativ zu W_0 wirklich ist, als auch A relativ zu W_0 wirklich ist. Dieser Satz, der wesentlich auf der Gültigkeit von (IV, 6) beruht, gilt jedoch nicht mehr im quanten-

logischen Modalkalkül. Dem entspricht der physikalische Sachverhalt, daß die Kenntnis des Zustandes $|f\rangle$ des Systems S die jeweils maximal mögliche Information darstellt, die man über S gewinnen kann, daß aber andererseits dieses Wissen nie ein vollständiges Wissen im Sinne von (IV, 16) ist. Eine Wirklichkeit W_0 , wie sie für alle klassischen Modalkalküle konstruiert werden kann, existiert also in der Quantenlogik nicht. Zu jedem Wissen W gibt es Aussagen A , die nicht objektiv in bezug auf W sind, die also nach dem oben Gesagten mit ihrem kontradiktorischen Gegenteil $\neg A$ im Modus der relativen Koexistenz sich befinden können.

Anhang

Anhang I. Quantentheorie

Es sollen hier einige für die vorliegende Arbeit wichtige Sätze und Begriffsbildungen angegeben werden, ohne daß im einzelnen auf die Begründung der Sätze eingegangen werden soll. Wir verweisen in diesem Zusammenhang auf die einschlägige Literatur [1], [10], [19], [20].

I. 1 Abgeschlossene Linearmannigfaltigkeiten

Es sei H der Hilbertsche Raum $|f_1\rangle, |f_2\rangle, \dots |g_1\rangle, |g_2\rangle \dots$ Elemente von H . Eine Teilmenge $\mathfrak{M}_A$ von H heißt eine lineare Mannigfaltigkeit, wenn sie mit den Elementen $|f_1\rangle, |f_2\rangle \dots |f_n\rangle$ auch deren Linearaggregate $\alpha_1 |f_1\rangle + \alpha_2 |f_2\rangle + \dots + \alpha_n |f_n\rangle$ enthält. ($\alpha_1, \dots, \alpha_n$ beliebige komplexe Zahlen.) Eine lineare Mannigfaltigkeit heißt abgeschlossen, wenn sie mit einer unendlichen Folge $\{|f_i\rangle\}$ von Elementen, falls diese einen Grenzwert $|f\rangle$ besitzt, auch dieses $|f\rangle$ enthält. Abgeschlossene Linearmannigfaltigkeiten wollen wir mit M_A, M_B, M_C bezeichnen. Ist $\mathfrak{N}_A$ irgendeine Menge von Elementen des Hilbertraumes, so ist die Menge $\{\beta_1 |g_1\rangle + \dots + \beta_n |g_n\rangle$ mit β_i beliebig komplex und $|g_i\rangle \in \mathfrak{N}_A$ stets eine lineare Mannigfaltigkeit, die wir als die von $\mathfrak{N}_A$ aufgespannte Linearmannigfaltigkeit bezeich-

nen. Fügen wir zu $\{\mathfrak{N}_A\} = \mathfrak{M}_A$ noch alle Häufungspunkte von $\{\mathfrak{N}_A\}$ hinzu, so entsteht die von $\mathfrak{N}_A$ aufgespannte abgeschlossene Linearmanigfaltigkeit $[\mathfrak{N}_A] = M_A$. Auf Grund dieser Definition ist

$$\mathfrak{N}_A \leq \{\mathfrak{N}_A\} \leq [\mathfrak{N}_A]$$

wobei $\{\mathfrak{N}_A\}$ und $[\mathfrak{N}_A]$ die kleinsten $\mathfrak{N}_A$ umfassenden Gebilde dieser Art sind.

Seien nun M_A, M_B abgeschlossene Linearmanigfaltigkeiten, $|f\rangle, |g\rangle$ deren Elemente, so bezeichnen wir die abgeschlossene Linearmanigfaltigkeit $[M_A, M_B]$ mit

$$M_A \cup M_B \Leftrightarrow [M_A, M_B]$$

$M_A \cup M_B$ besteht offenbar aus allen Elementen $\alpha|f\rangle + \beta|g\rangle$ und deren Häufungspunkten.

Ist M_A Teilmenge von M_B , also $M_A \leq M_B$, so betrachten wir die Gesamtheit der Elemente von M_B , die orthogonal zu allen Elementen von M_A sind. Auch dies ist eine abgeschlossene Linearmanigfaltigkeit, die wir mit $M_B - M_A$ bezeichnen wollen. Insbesondere heißt $H - M_A$ die zu M_A komplementäre abgeschlossene Linearmanigfaltigkeit.

Weiter wollen wir mit $M_A \cap M_B$ diejenige abgeschlossene Linearmanigfaltigkeit bezeichnen, die aus den gemeinsamen Elementen von M_A und M_B besteht. $M_A \cap M_B$ werden wir auch als Durchschnitt von M_A und M_B bezeichnen. Schließlich erwähnen wir drei besonders einfache abgeschlossene Linearmanigfaltigkeiten. Der Hilbertraum H , die Nullmenge 0 und die Menge aller $a|f\rangle$ (a variabel).

I. 2 Projektionsoperatoren

Der Begriff der Projektion kann hier auf Grund des folgenden Satzes eingeführt werden: Ist M_A eine abgeschlossene Linearmanigfaltigkeit, dann kann ein beliebiges $|f\rangle \leq H$ auf eine und nur eine Weise in eine Summe $|f\rangle = |f_A\rangle + |f_{\bar{A}}\rangle$ mit $|f_A\rangle \leq M_A$, $|f_{\bar{A}}\rangle \leq H - M_A$ zerlegt werden. $|f_A\rangle$ nennen wir die Projektion von $|f\rangle$ in M_A und schreiben $P_A|f\rangle = |f_A\rangle$. P_A

ist also eine Operation, die jedem $|f\rangle$ seine Orthogonalprojektion $|f_A\rangle$ in M_A zuordnet. P_A ist somit ein überall in H definierter Operator, der als Projektionsoperator bezeichnet werde.

Satz 1: Ein Projektionsoperator P_A hat die folgenden Eigenschaften:

- a) linear $P_A(\alpha_1|f_1\rangle + \alpha_2|f_2\rangle) = \alpha_1 P_A|f_1\rangle + \alpha_2 P_A|f_2\rangle$
- b) hermitisch $P_A^+ = P_A$
- c) idempotent $P = P^2$
- d) M_A ist die Menge aller Werte von P_A , also die Menge aller $P_A|f\rangle$
- e) M_A kann auch durch die Menge der Lösungen von $P_A|f\rangle = |f\rangle$ gekennzeichnet werden. $H - M_A$ ist dann die Menge aller Lösungen von $P_A|f\rangle = 0$.

Ein Projektionsoperator hat somit die Eigenwerte 0 und 1, wobei der Eigenwert 0 zu allen $|f\rangle \in H - M_A$, der Eigenwert 1 zu allen Elementen $|f\rangle \in M_A$ gehört. Unabhängig von den abgeschlossenen Linearmanigfaltigkeiten kann man die Projektionsoperatoren auch auf Grund des folgenden Satzes definieren.

Satz 2: Ein überall sinnvoller Operator P ist dann und nur dann ein Projektionsoperator, wenn er die Eigenschaften

$$P^2 = P \qquad P^+ = P$$

besitzt.

Das Problem, welche Zusammensetzungen von Projektionsoperatoren wieder Projektionsoperatoren sind, wird durch den folgenden Satz beantwortet:

Satz 3: Sind P_1 und P_2 Projektionsoperatoren, die auf M_1 und M_2 projizieren, so ist

- a) $P_1 \cdot P_2$ dann und nur dann ein Projektionsoperator, wenn $P_1 \cdot P_2 = P_2 \cdot P_1$ ist. $P_1 \cdot P_2$ projiziert dann auf $M_1 \cap M_2$.
- b) $P_1 + P_2 - P_1 P_2$ ist dann und nur dann ein Projektionsoperator, wenn $P_1 \cdot P_2 = P_2 \cdot P_1$ ist. $P_1 + P_2 - P_1 P_2$ projiziert dann auf $M_1 \cup M_2$.

c) $P_1 - P_2$ ist dann und nur dann ein Projektionsoperator, wenn $P_1 \cdot P_2 = P_1$ ist. $P_1 - P_2$ projiziert dann auf $M_1 \cup (H - M_2)$.

Insbesondere ist also $1 - P_A$ derjenige Operator, der auf $H - M_A$ projiziert, da wegen $P_H = 1$ stets $P \cdot P_H = P$ gilt. Weiterhin gilt der Satz, daß M_1 und M_2 dann und nur dann orthogonal sind, wenn $P_1 \cdot P_2 = 0$ ist. Es ist dann also $M_1 \cap M_2 = 0$.

Satz 4: Sind P_1 und P_2 Projektionsoperatoren, die auf M_1 und M_2 projizieren, so sind die folgenden Behauptungen äquivalent:

- a) $P_1 \leq P_2$
- b) $\|P_1|f\rangle\| \leq \|P_2|f\rangle\|$
- c) $P_1 = P_1 P_2 = P_2 P_1$
- d) $M_1 \leq M_2$

Satz 5: Sind $P_1, P_2 \dots P_N$ Projektionsoperatoren, $M_1, M_2 \dots M_N$ die entsprechenden Linearmannigfaltigkeiten, so ist $P = P_1 + P_2 \dots + P_N$ dann und nur dann ein Projektionsoperator, wenn $P_i \cdot P_k = 0$ ist, für alle $i \neq k$. P projiziert dann auf $M = M_1 \cup M_2 \dots \cup M_N$.

Satz 6: Ist $P_1 \dots P_N$ eine aufsteigende Folge von Projektionsoperatoren, $P_1 \leq P_2 \dots$, so konvergiert diese gegen einen Projektionsoperator P in dem Sinne, daß für alle $|f\rangle : P_n |f\rangle \rightarrow P |f\rangle$. Es sind dann alle $P_n \leq P$.

Hat man also eine Folge $P_1 \dots P_n \dots$ von paarweise vertauschbaren Projektionsoperatoren, so sind (Satz 5) $P_1, P_1 + P_2, P_1 + P_2 + P_3$ Projektionsoperatoren und offenbar ansteigend, konvergieren also (Satz 6) gegen ein P , welches größer ist als sie alle. $P = \sum_1^\infty P_i$ projiziert dann auf $M = M_1 \cup M_2 \dots$.

I.3 Quantenmechanische Eigenschaften

Wir wollen hier diejenige Formulierung der Quantenmechanik verwenden, die sich des Begriffes der „Eigenschaft“ eines Systems S bedient. Eigenschaften sind z. B., daß gewisse Obser-

vale von S bestimmte Meßwerte besitzen, oder daß die Meßwerte in einem vorgegebenen Wertebereich liegen usw. Um alle diese Eigenschaften in einer einheitlichen Weise zu schreiben, führen wir die folgende Definition ein: Liegt der Zustand $|f\rangle$ eines vorgegebenen Systems S in einer abgeschlossenen Linearmannigfaltigkeit M_A (die natürlich auch aus einem einzigen Zustand $|A\rangle$ bestehen kann), so wollen wir sagen, daß das System S die Eigenschaft E_A hat. Nach dem oben über Projektionsoperatoren Gesagten ist das genau dann der Fall, wenn

$$P_A |f\rangle = |f\rangle$$

gilt, wobei P_A auf M_A projiziert. Da P_A ein linearer, selbstadjungierter Operator mit den Eigenwerten 1 und 0 ist, kann man P_A als den Operator einer Observablen auffassen, die die Meßwerte 1 und 0 hat. Das Vorliegen des Eigenwertes 1 besagt dann, daß das im Zustande $|f\rangle$ befindliche System S die Eigenschaft E_A hat. Bei Vorliegen des Eigenwertes 0 liegt $|f\rangle$ in $M'_A = H - M_A$, so daß S die Eigenschaft E'_A besitzt.

Ebenso wie bei beliebigen Observablen sind daher auch zwei oder mehrere Eigenschaften E_A, E_B, E_C an einem System genau dann gleichzeitig meßbar, wenn die zugehörigen Projektionsoperatoren miteinander vertauschbar sind. Die Kommensurabilität von zwei Eigenschaften E_A und E_B stellt dabei eine Verschärfung des an beliebigen Observablen definierten Begriffes der gleichzeitigen Meßbarkeit dar. Denn z. B. ist die Eigenschaft E_1 , daß die Observable R in dem im Zustand $|f\rangle$ befindlichen System S den Wert R_1 hat und die Eigenschaft E_2 , daß T den Wert T_2 hat (R und T mögen ein diskretes, nicht entartetes Spektrum besitzen), gleichzeitig meßbar, wenn $P_1 = |R_1\rangle\langle R_1|$ und $P_2 = |T_2\rangle\langle T_2|$ vertauschbar sind, während für die Kommensurabilität von R und T die Vertauschbarkeit aller $|R_j\rangle\langle R_j|$ mit allen $|T_k\rangle\langle T_k|$ erforderlich ist.

Häufig wird auch statt gleichzeitiger Meßbarkeit bei Eigenschaften der Begriff der gleichzeitigen Entscheidbarkeit verwendet. [1] Wir haben diesen Begriff hier absichtlich vermieden, da seine Verwendung erst dann sinnvoll ist, wenn die Negation einer Aussage eingeführt worden ist, was erst in Kap. II mit rein logischen Mitteln geschehen soll.

Anhang II. Verbandstheoretische Hilfsmittel

Es sollen im folgenden Abschnitt einige für die vorliegende Arbeit wichtige verbandstheoretische Begriffe und Sätze zusammengestellt werden. Es soll hier jedoch nicht in irgendeiner Hinsicht eine Vollständigkeit des dargebotenen Materials angestrebt werden. Wegen aller Einzelheiten, Beweise usw. sei daher auf die ausführliche Darstellung dieser Fragen bei Birkhoff [21] und Hermes [22] verwiesen.

Eine Menge von Elementen A, B, C heißt eine Halbordnung oder eine halbgeordnete Menge, wenn zwischen je zwei Elementen eine Relation $A \leq B$ erklärt ist, so daß gilt:

$$(A \text{ II, } 1) \quad A \leq A$$

$$(A \text{ II, } 2) \quad A \leq B, \quad B \leq C \Rightarrow A \leq C$$

$$(A \text{ II, } 3) \quad A \leq G, \quad B \leq A \Rightarrow A = B$$

Zwei Elemente A und B sollen vergleichbar heißen, wenn zwischen ihnen wenigstens eine der Relationen $A \leq B$ oder $B \leq A$ gilt. Eine Halbordnung, in der je zwei Elemente vergleichbar sind, heißt eine Kette.

Ein Element, welches kein anderes Element umfaßt, heiße minimal, und ein Element, das von keinem umfaßt wird, maximal. Ein Element, welches in jedem anderen enthalten ist, heiße kleinstes oder Nullelement $\wedge$, und ein Element, welches alle anderen umfaßt, größtes oder Einselement $\vee$.

Ist weiter $\mathfrak{A}$ eine Teilmenge einer Halbordnung, so heißt ein Element S obere bzw. untere Schranke von $\mathfrak{A}$, wenn für alle $a \in \mathfrak{A}$ $a \leq S$ bzw. $S \leq a$ gilt. Existiert eine solche obere bzw. untere Schranke, so heißt $\mathfrak{A}$ nach oben bzw. nach unten beschränkt. Die kleinste obere Schranke heißt obere Grenze, entsprechend die größte untere Schranke untere Grenze.

Gibt es nun in einer Halbordnung zu je zwei Elementen A, B ein Element $A \wedge B$, welches untere Grenze von A und B ist, also

$$(A \text{ II, } 4) \quad A \wedge B \leq A$$

$$(A \text{ II, } 5) \quad A \wedge B \leq B$$

$$(A \text{ II, } 6) \quad C \leq A, \quad C \leq B \Rightarrow C \leq A \wedge B$$

erfüllt, so heißt die Menge ein Halbverband. Existiert darüber hinaus für je zwei Elemente auch stets eine obere Grenze $A \vee B$, die also

$$(A \text{ II, } 7) \quad A \leq A \vee B$$

$$(A \text{ II, } 8) \quad B \leq A \vee B$$

$$(A \text{ II, } 9) \quad A \leq C, B \leq C \Rightarrow A \vee B \leq C$$

erfüllt, so heißt die untersuchte Menge ein Verband. In Verbänden gilt das Dualitätsprinzip, welches besagt, daß die duale Aussage eines jeden Satzes wieder ein Satz ist. Man geht von einer Aussage zu der dualen über, indem man $\leq$ durch $\geq$ bzw. $\geq$ durch $\leq$ ersetzt, und $\wedge$ durch $\vee$ bzw. $\vee$ durch $\wedge$. Ein Beispiel dafür sind die stets in Verbänden gültigen Beziehungen:

$$(A \text{ II, } 10) \quad A \vee (B \wedge C) \leq (A \vee B) \wedge (A \vee C)$$

$$(A \text{ II, } 11) \quad (A \wedge B) \vee (A \wedge C) \leq A \wedge (B \vee C)$$

In einem Verband existiert auch für jede endliche Menge von Elementen $\{A_1, \dots, A_n\}$ eine obere und eine untere Grenze. Ist dies sogar für jede unendliche Menge $\{A_1, \dots\}$ von Elementen der Fall, so heißt der betreffende Verband vollständig. Es gilt dann der Satz, daß jeder vollständige Verband ein Null- und Einselement besitzt.

Gibt es in einem Verband zu je zwei Elementen A und B ein weiteres Element $B \rightarrow A$ mit der Eigenschaft

$$(A \text{ II, } 12) \quad A \wedge B \rightarrow A \leq B$$

$$(A \text{ II, } 13) \quad A \wedge C \leq B \Rightarrow C \leq B \rightarrow A$$

so daß $B \rightarrow A$ also das größte unter den Elementen $X_1, X_2, \dots$ ist, welche die Bedingung $A \wedge X_i \leq B$ erfüllen, so heißt der betreffende Verband relativ-pseudokomplementär. $B \rightarrow A$ ist dann das zu B relative Pseudokomplement von A . Existiert weiter ein Nullelement $\wedge$, so heißt $\wedge \rightarrow A$ das Pseudokomplement zu A , und möge mit $\neg A$ bezeichnet werden. Es erfüllt dann die beiden Beziehungen

$$(A \text{ II, } 14) \quad A \wedge \neg A \leq \wedge$$

$$(A \text{ II, } 15) \quad A \wedge C \leq \wedge \Rightarrow C \leq \neg A$$

In relativ-pseudokomplementären Verbänden existiert weiter stets ein Einselement, nämlich $V = A \rightarrow A$, wobei A ein beliebiges Element ist. Eine weitere sehr wichtige Eigenschaft dieser Verbände ist, daß sie distributiv sind, d. h. es gilt über (A II, 10, 11) hinaus

$$(A \text{ II, } 16) \quad A \vee (B \wedge C) = (A \vee B) \wedge (A \vee C)$$

$$(A \text{ II, } 17) \quad (A \wedge B) \vee (A \wedge C) = A \wedge (B \vee C)$$

Ein beliebiger Verband heißt relativ-komplementär, wenn für je zwei Elemente A und B , die die Bedingung $A \leq B$ erfüllen, folgendes gilt: Zu jedem Element X existiert ein Element Y , das zu A und B relative Komplement von X , welches die beiden Bedingungen

$$(A \text{ II, } 18) \quad X \wedge Y = A \quad X \vee Y = B$$

erfüllt. Existiert weiter ein Nullelement $\wedge$, so daß $\wedge \leq B$ für alle B gilt, so heißt der Verband abschnittskomplementär. Ist außerdem ein Element $\vee$ vorhanden, so heißt der Verband komplementär. Zu jedem X gibt es dann ein Y , das Komplement zu X , so daß also

$$(A \text{ II, } 19) \quad X \wedge Y = \wedge \quad X \vee Y = \vee$$

gilt. Im allgemeinen ist die Bildung der Komplemente nicht eindeutig. Es gilt jedoch der wichtige Satz, daß in einem distributiven Verband ein Element höchstens ein Komplement besitzt. Existiert also in einem distributiven Verband ein Komplement, so ist es dann eindeutig bestimmt. Ein komplementärer distributiver Verband heißt auch ein Boolescher Verband oder auch eine Boolesche Algebra. Verbände, in denen die Distributivgesetze (A II, 16, 17) nur unter den Bedingungen $A \leq C$ bzw. $C \leq A$ gelten, heißen modular. Ist ein modularer Verband komplementär, so ist er auch abschnittskomplementär; ist er abschnittskomplementär, so auch relativkomplementär.

Im Gegensatz zu distributiven ist jedoch in modularen Verbänden die Komplementbildung nicht mehr eindeutig. Es wird im allgemeinen zu einem Element A mehrere Komplemente $A, A' \dots$ geben, die den Bedingungen

$$(A \text{ II, } 20) \quad A \wedge A' = \wedge \quad A \vee A' = \vee$$

genügen. In gewissen Fällen läßt sich jedoch ein eindeutig definiertes Komplement angeben. Dies ist der Fall, wenn der untersuchte Verband orthokomplementär ist. Ein Verband heißt orthokomplementär dann und nur dann, wenn er einen dualen Automorphismus $A \Rightarrow A'$ gestattet, der außer (A II, 20) noch der Bedingung

$$(A \text{ II, 21}) \quad (A')' = A$$

genügt. (Eine Abbildung $X \Rightarrow X'$ eines Verbandes auf sich selbst heißt ein dualer Automorphismus, wenn er zwei Elemente A, B , die in der Relation $A \leq B$ stehen, in zwei Elemente A, B abbildet, die in der dualen Relation $B \leq A$ stehen). Dieses (eindeutig definierte) Komplement wird als das Orthokomplement zu A bezeichnet.

Anhang III. Operative Logik

Im folgenden Abschnitt soll kurz skizziert werden, wie sich die in den mathematischen Theorien gebräuchliche Logik – abgesehen von dem tertium non datur – operativ begründen läßt. Es ist hier nicht möglich, diesen Aufbau in allen Einzelheiten darzustellen, weshalb wir uns oft auf kurze Hinweise beschränken müssen. Im übrigen sei auf die ausführliche Darstellung dieses Problems bei Lorenzen [9] verwiesen.

Wir wollen ausgehen von einem Kalkül K_0 , also von einigen Figuren $X, Y, Z, \dots$ und Regeln der Form $x \rightarrow y$, die besagen, daß, wenn eine Figur x schon hergestellt ist, dann auch y herzustellen sei. „ x “ steht hier als eine Variable für Figuren, die selbst schon in K_0 hergestellt sind. Wir definieren dann einen weiteren Kalkül K , der aus den Figuren $A, B, C, \dots$, den Anfängen des Kalküls, und Regeln wie $a \rightarrow b$; $a, b \rightarrow x$ bestehen möge. „ a “ „ b “ stehen jetzt als Variable für alle „Aussagen“ des Kalküls. Unter Aussage wollen wir hierbei eine Figur verstehen, die in K_0 konstruierbar ist. K_0 dient also nur als Hilfskalkül, um festzulegen, was unter einer Aussage in K zu verstehen ist. „ $a \rightarrow b$ “ bedeutet wieder, daß, wenn a hergestellt ist, dann auch b herzustellen sei. $a, b \rightarrow c$ bedeutet, daß, wenn a und b schon hergestellt seien, dann c herzustellen ist. Eine mit den Anfängen und Regeln

von K herstellbare Figur heie eine „ableitbare Aussage“. Unter einer Ableitung α wollen wir im folgenden ein aus mehreren Zeilen bestehendes Schema verstehen, in dem alle zur Herstellung einer ableitbaren Aussage notwendigen Schritte explizit aufgefhrt werden. Die letzte Zeile einer Ableitung α werde als „Endaussage“ $\varepsilon(\alpha)$ bezeichnet. Das Zeichen $A \vdash_K B$ soll weiter heien, da man nach Hinzunahme von A zu den Anfngen des Kalkls K die Figur B nach den Regeln von K herstellen kann. Wenn keine Verwechslungsmglichkeiten bestehen, wollen wir auch einfach $\vdash$ schreiben.¹

Eine Regel R heie zulssig in einem Kalkl K , wenn man jede Ableitung α , die die Regel R bentzt, ersetzen kann durch eine Ableitung α' , die R nicht bentzt, aber die die gleiche Endaussage $\varepsilon(\alpha) = \varepsilon(\alpha')$ besitzt. Eine solche Umformung der Ableitung α in α' wollen wir ein Eliminationsverfahren nennen. Wir sagen auch, da die Benutzung der Regel R in jeder Ableitung α von K eliminiert werden kann. Die Elimination eines „Anfangs“ ist in dieser Definition inbegriffen, wenn wir die Anfnge als spezielle Regeln ohne Vorderformeln betrachten.

Die verschiedenen Verfahren, nach denen eine Regel eliminiert werden kann, werden in der Protologik untersucht. Fr den Gebrauch der operativen Mathematik gengt es, sich auf die fnf protologischen Prinzipien zu beschrnken: Deduktionsprinzip, Induktionsprinzip, Inversionsprinzip, Gleichheitsprinzip und Unableitungsprinzip. Diese Prinzipien, auf die wir im einzelnen hier nicht eingehen knnen, geben verschiedene allgemeine Verfahren an, nach denen man Regeln, die man zur Herstellung von Figuren benutzt, in einer solchen Ableitung eliminieren kann. Es ist wichtig, im Auge zu behalten, da alle in der Protologik diskutierten Eliminationsverfahren von der Vorstellung ausgehen, da es sich dabei um eine materielle Konstruktion von Figuren (z. B. aus Steinchen) handelt, deren Herstellung keinen anderen als den durch die Regeln des Kalkls festgelegten Einschrnkungen unterworfen ist.

¹ Eine Aussage A heie weiter unableitbar, wenn A ungleich allen ableitbaren Aussagen ist. Um zu przisieren, was gleich und ungleich bei allgemeinen Figuren heit, hat man auch die Gleichheit und die Ungleichheit durch einen Kalkl zu definieren. Wir wollen hierauf jedoch nicht nher eingehen.

Außer den Kalkülen kann man noch sog. „Metakalküle“ untersuchen, deren „Aussagen“ die Regeln der oben untersuchten Kalküle K sind. Weiter bezeichnen wir als eine Metaregel eine Regel, die die Hinzunahme einer weiteren Regel zum Metakalkül gestattet, und die allgemeine Form $A \rightarrow B \dot{\rightarrow} C \rightarrow D$ haben wird. Der einfache Pfeil „ $\rightarrow$ “ ist jetzt als eine bedeutungsfreie Figur anzusehen, während „ $\dot{\rightarrow}$ “ jetzt genau die Bedeutung hat, die „ $\rightarrow$ “ früher in Kalkülen hatte. Entsprechend wollen wir von einer „zulässigen“ Metaregel sprechen, wenn sie eine Regel über Regeln von K ist, die angewandt auf zulässige Regeln von K stets wieder zulässige Regeln liefert.

Ebenso kann man fortschreitende Metakalküle usw. untersuchen. Als gemeinsamen Namen für Aussagen, Regeln, Metaregeln gebrauchen wir jetzt den Terminus „Aussage“ und benutzen A, B, C als Variable für Aussagen. Wir wollen nun fragen ob es Regeln, Metaregeln usw. gibt, die in jedem Kalkül zulässig sind. Solche Aussagen wollen wir als allgemein zulässige Aussagen bezeichnen. Ohne auf die Beweise einzugehen, sei bemerkt, daß sich zunächst die folgenden Regeln zur Ableitung von allgemein zulässigen Aussagen aufstellen lassen:

- | | |
|------------|--|
| (A III, 1) | $A \rightarrow A$ |
| (A III, 2) | $A \rightarrow B, B \rightarrow C \Rightarrow A \rightarrow C$ |
| (A III, 3) | $A, C \dot{\rightarrow} B \Leftrightarrow C \dot{\rightarrow} A \rightarrow B$ |

Während „ $\rightarrow$ “ Bestandteil der Aussagen ist, haben wir hier $\Rightarrow$ zur Mitteilung der Regeln verwendet. $X \Rightarrow Y$ soll also heißen: Ist X ableitbar, so auch Y . Der Doppelpfeil $\Leftrightarrow$ steht an Stelle von $\Rightarrow$ und $\Leftarrow$, so daß (A III, 3) also zwei Regeln sind. Das Komma „ $,$ “ zwischen zwei Aussagen A, B wurde hier in demselben Sinne wie oben gebraucht, daß also sowohl die Aussage A als auch die Aussage B abgeleitet ist.

Die Gesamtheit der allgemein zulässigen Regeln, Metaregeln usw. bezeichnen wir als operative Logik. Die Regeln der Logik sind also dadurch gekennzeichnet, daß man mit ihrer Hilfe keine Aussage, Regel usw. herleiten kann, die nicht auch schon ohne Verwendung der logischen Regeln abgeleitet werden könnte.

Zur Vereinfachung des Ableitens in Kalkülen ist es üblich, über die durch (A III, 1, 2, 3) gegebenen Regeln hinaus einige weitere Symbole durch relativ zulässige Regeln einzuführen. Ein Beispiel dafür ist die Regel, durch die das Zeichen $\wedge$, die Konjunktion, eingeführt wird:

$$(A \text{ III}, 4) \qquad A, B \rightarrow A \wedge B$$

Eine solche Regel ist sicherlich nicht zulässig, denn die Aussage $A \wedge B$ ist jetzt u. U. ableitbar, während dies vorher sicher nicht der Fall war, da angenommen werden soll, daß $\wedge$ bisher nicht unter dem Zeichen des Kalküls vorgekommen ist. (A III, 4) ist hingegen relativ zulässig. Damit ist gemeint, daß aus jeder Ableitung α , die eine Endaussage $\varepsilon(\alpha)$ besitzt, in der $\wedge$ nicht vorkommt, die Benutzung der Regel (A III, 4) eliminiert werden kann.

Die Konjunktion, Disjunktion und Partikularisation werden auf diese Weise in die Logik eingeführt. Wir wollen an dieser Stelle nicht näher darauf eingehen, da der Aufbau der operativen Logik in Kap. II an dem hier besonders wichtigen Beispiel des Kalküls der kommensurablen Eigenschaften im einzelnen durchgeführt werden soll. Auch die Begriffe der Negation und der Generalisation sollen dort ausführlich behandelt werden, weshalb hier auf eine Diskussion dieser Begriffe verzichtet werden soll.

Literatur

- [1] J. v. Neumann: Math. Grundl. der Quantenmechanik. Berlin 1932.
- [2] G. Birkhoff, J. v. Neumann: Ann. of Math. 37, 823 (1936).
- [3] D. Hilbert, P. Bernays: Grundlagen d. Mathematik. Berlin 1934/39.
- [4] D. Hilbert, W. Ackermann: Grundzüge d. Theor. Logik. Berlin 1949.
- [5] Th. Skolem: Skrifter utgit av. Videnskapsselskapet Kristiana 6 (1923).
- [6] H. B. Curry: Outlines of a formalist philosophy of Mathematics. Amsterdam 1952.
- [7] P. Lorenzen: Math. Z. 53 (1950), 54 (1951).
- [8] P. Lorenzen: Formale Logik. Berlin 1959.

- [9] P. Lorenzen: Einführung in die operative Logik und Mathematik. Berlin 1955.
 - [10] G. Süßmann: Abh. d. Bayer. Akad. der Wissensch. 88 (1958).
 - [11] C. F. v. Weizsäcker: Naturwissenschaften 42, 521 (1955).
 - [12] C. F. v. Weizsäcker: Naturwissenschaften 42, 545 (1955).
 - [13] C. F. v. Weizsäcker: Z. f. Naturforschg. 13a, 245, 704 (1958).
 - [14] H. Reichenbach: Philosophische Grundlagen der Quantenmechanik. Basel 1949.
 - [15] J. P. Destouches-Fevrier: La structure des théories physiques. Paris 1951.
 - [16] H. B. Curry: Leçons de Logique Algébrique. Paris-Louvain 1952.
 - [17] O. Becker: Untersuchungen über den Modalkalkül. Meisenheim 1952.
 - [18] C. J. Lewis: Langford, Symbolic Logic. New York 1932.
 - [19] G. Ludwig: Grundlagen der Quantenmechanik. Berlin 1954.
 - [20] P. R. Halmos: Introduction to Hilbert-Space. New York 1951.
 - [21] G. Birkhoff: Lattice Theorie. New York 1948.
 - [22] H. Hermes: Einführung in die Verbandstheorie. Berlin 1955.
-