

Sitzungsberichte

der

mathematisch-naturwissenschaftlichen
Abteilung

der

Bayerischen Akademie der Wissenschaften

zu München

1925. Heft II

Juli- bis Dezembersitzung

München 1925

Verlag der Bayerischen Akademie der Wissenschaften
in Kommission des G. Franz'schen Verlags (J. Roth)



Die Kausalstruktur der Welt und der Unterschied von Vergangenheit und Zukunft.

Von Hans Reichenbach.

Vorgelegt von C. Carathéodory in der Sitzung am 7. November 1925.

I. Der Determinismus und das Problem des „jetzt“.

Es ist üblich geworden, die Kausalhypothese der Physik als eine so selbstverständliche Notwendigkeit zu betrachten, daß man nicht mehr daran denkt, sie einer Kritik zu unterziehen. Dabei bemerkt man meist gar nicht, in wie hohem Grade diese Hypothese eine Extrapolation über den erfahrungsgemäßen Tatbestand hinaus bedeutet; in der Annahme, daß ohne sie keine exakte Naturerkenntnis möglich sei, erschöpft sich die gewöhnliche Verteidigung dieses Standpunkts. Im folgenden soll gezeigt werden, daß auch ohne die Hypothese einer strengen Kausalität eine quantitative Beschreibung des Naturgeschehens möglich ist, die gerade alles das leistet, was die Physik überhaupt leisten kann, und die überdies noch geeignet ist, die Frage nach dem Unterschied von Vergangenheit und Zukunft zu lösen, auf welche die strenge Kausalhypothese keine Antwort hat.

Wir müssen dieser Untersuchung eine Unterscheidung vorausschicken, die allein schon das Problematische der Kausalhypothese hervortreten läßt. Die erste Form der Kausalhypothese liegt vor, wenn die Physik Gesetze aufstellt, d. h. Aussagen macht von der Form: „wenn A ist, dann ist B “. Wir nennen sie die Implikationsform der Kausalhypothese. Die zweite Form aber geht darüber hinaus und behauptet etwas über den Ablauf der Welt als Ganzes; sie besagt nämlich, daß dieser Ablauf unveränderlich feststehe, daß mit einem einzigen Querschnitt

der vierdimensionalen Welt Vergangenheit und Zukunft völlig bestimmt seien. Diese Behauptung, die man auch Determinismus nennt, wollen wir die Determinationsform der Kausalhypothese nennen. Es ist offensichtlich, daß die zweite Behauptung sehr viel weiter geht als die erstere; und es erscheint außerordentlich kühn, daß die Naturwissenschaft den Schritt von der immerhin noch plausiblen Implikationsform zu diesem Anspruch auf Beherrschung des Weltablaufs gemacht hat. Man rechtfertigt ihn, indem man die beiden Formen in einen Zusammenhang bringt; aber man bemerkt nicht, daß dabei zu der Implikationsform der Kausalhypothese noch eine zweite Annahme hinzutritt, die zugleich gegenüber dem Erfahrungsmaterial eine sehr zweifelhafte Behauptung darstellt. Diese zusätzliche Hypothese läßt sich erkennen, wenn man den Übergang von der Implikationsform zur Determinationsform genauer betrachtet.

Wenn die Implikationsform besagt, daß die Ursache A mit Gewißheit die Wirkung B hat, so stellt sie diese Behauptung doch nur für den Fall auf, daß die Ursache A in aller Strenge wirklich vorliegt. Aber gerade dies ist bekanntlich nie erfüllt, so daß bei jeder Anwendung der Implikationsform auf die Wirklichkeit noch eine zweite Hypothese notwendig wird, die sich auf den Rest von Faktoren bezieht, welche außer A noch da sind. Man formuliert diese zusätzliche Hypothese gewöhnlich als die Annahme, daß die Restfaktoren nur einen quantitativ kleinen Einfluß haben. Aber das ist nicht genau. Die Annahme lautet in Wahrheit, daß die Restfaktoren nach den Gesetzen der Wahrscheinlichkeitsrechnung ihren Einfluß ausüben. Es können gelegentlich wohl größere Störungen vorkommen, aber bei wiederholten Fällen entsprechen die Störungen einem statistischen Gesetz. Wie ich an anderer Stelle gezeigt habe,¹⁾ ist dies die Annahme, daß die Störungen durch eine stetige Wahrscheinlichkeitsfunktion geregelt sind. Dieses Wahrscheinlichkeitsprinzip tritt stets hinzu, wenn die Kausalhypothese in ihrer Im-

¹⁾ Der Begriff der Wahrscheinlichkeit für die mathematische Darstellung der Wirklichkeit. Diss. Erlangen 1915 und Zschr. f. Philos. u. philos. Kritik 161, 1917, S. 209. Vgl. auch Naturwiss. 1920, S. 46 und 146. — Die stetige Funktion ist nicht immer die Gaußsche, diese gilt nur für besondere Fälle.

plikationsform auf die Wirklichkeit angewandt wird. Es läßt sich nicht etwa aus der Implikationsform ableiten, sondern bedeutet eine selbständige Annahme, ohne welche die Implikationsform wertlos wäre; denn man könnte sie sonst in keinem Fall auf die Wirklichkeit anwenden. Die physikalische Erkenntnis beruht deshalb auf zwei Prinzipien, dem Prinzip der kausalen Verknüpfung und dem Prinzip der wahrscheinlichkeitsgemäßen Verteilung.

Wie kommen wir nun von hier aus zur Determinationsform? Um diesen Zusammenhang aufzudecken, wollen wir die Determinationsform der Kausalhypothese für eine mit stetigem materiellem Feld erfüllte Welt formulieren. In einer solchen Welt brauchen wir dann nicht von einzelnen Ereignissen zu reden, sondern können die Welt durch Angabe der Feldverteilung völlig beschreiben.

Die Determinationshypothese lautet dann: Ist für einen Querschnitt $t = \text{konst.}$ der vierdimensionalen Welt die Feldverteilung und außerdem die ersten und zweiten Ableitungen der Feldgrößen nach der Zeit gegeben, so ist Vergangenheit und Zukunft völlig bestimmt. Dabei kann man sich die Feldverteilung so vorstellen, daß etwa der Einsteinsche Tensor T_{ik} als Funktion gewählter Raumkoordinaten in aller Strenge gegeben ist.

Vergleichen wir diese Aussage mit dem Erfahrungsmaterial, so finden wir eben den Unterschied, auf den wir bereits hingewiesen haben. Einerseits ist der Feldzustand nie mit völliger Strenge gegeben, andererseits werden die daraus berechneten früheren und späteren Zustände stets nur mit Wahrscheinlichkeit festgelegt. Von der Erfahrung aus läßt sich also die Determinationshypothese nur durch einen Grenzübergang gewinnen, der die approximative Feldverteilung in die strenge und die Wahrscheinlichkeit in Gewißheit verwandelt. Das Problematische dieses Grenzübergangs ist es, was mit der unkritischen Aufstellung des Determinismus zumeist übersehen wird.

Denken wir etwa die Verteilung der Materie innerhalb der Erdkugel durch ihre Dichte σ als Funktion der Koordinaten gegeben. Für die Astronomie wird der Ansatz $\sigma = \text{konst.}$ genügen. Für die Geologie wird σ entsprechend den Erdschichtungen variabel sein. Die Physik geht sehr viel weiter und will die Lage jedes einzelnen Moleküls kennen, d. h. eine Dichtefunktion gewinnen,

die sehr viel feinere räumliche Schwankungen macht als die geologische Dichte. Jede dieser Genauigkeitsstufen liefert eine Vorausbestimmung des Geschehens, d. h. zukünftiger Dichtezustände; mit wachsender Genauigkeit steigt die Sicherheit des berechneten Resultats. Die Determinationshypothese nimmt nun an, daß es eine Funktion gibt (ev. unter Aufspaltung des Skalars σ in einen Tensor T_{ik}), welche das Resultat mit Gewißheit bestimmt.

Sei es zugegeben, daß der Grad der Wahrscheinlichkeit beliebig nahe an 1 gesteigert werden kann — so bleibt in der Determinationshypothese doch die Annahme enthalten, daß die Reihe der Feldfunktionen, in wachsender Genauigkeit geordnet, eine Grenze hat. Im Sinne des Erfahrungsmaterials liegt nur die Aussage, daß zu jeder Feldfunktion eine genauere existiert, welche eine höhere Wahrscheinlichkeit liefert. Es geht weit darüber hinaus, zu sagen, daß es in dieser Reihe eine letzte Funktion gibt, die dann die Wahrscheinlichkeit 1 liefert. Dies ist die Extrapolation, die in der Determinationshypothese enthalten ist.

Man kann natürlich nicht ohne weiteres sagen, daß diese Extrapolation falsch ist; aber man kann behaupten, daß alles, was mit dieser Extrapolation erklärbar ist, auch ohne sie erklärt werden kann. Denn für alle kontrollierbaren physikalischen Aussagen wird stets nur die Tatsache der nach 1 steigerungsfähigen Genauigkeit benutzt, niemals die Existenz der Grenzfunktion selbst. Die Determinationshypothese ist deshalb für die Physik völlig leer; und wenn man sie auch nicht direkt widerlegen kann, so gibt es doch jedenfalls nichts, was für sie spricht. Im folgenden soll deshalb diese Hypothese weggelassen und gezeigt werden, wie sich die Kausalstruktur der Welt allein mit Hilfe des Begriffs der wahrscheinlichen Bestimmtheit beherrschen läßt.

Man hat den Wert der Determinationshypothese darin gesehen, daß sie den Wahrscheinlichkeitsbegriff in der Naturerklärung eliminiert. Für sie ist das Wahrscheinlichkeitsprinzip nur ein Aushilfsmittel, das man benutzt, so lange man die genauen Bedingungen eines Vorgangs nicht kennt; bei völlig genauer Kenntnis aller Umstände würde dieses Aushilfsmittel überflüssig. Aber diese Rechtfertigung vergift, daß der Wahrscheinlichkeitsbegriff eben nur für die Grenze wegfällt; für jede praktisch mögliche Aussage der Naturwissenschaft ist er dann immer noch

unentbehrlich. Denn wenn die Grenzfunktion auch existiert, so ist sie doch niemals in aller Strenge bekannt. Aber die Tatsache, daß für die ungenauen Beschreibungen des Naturgeschehens nun wenigstens Wahrscheinlichkeitsgesetze gelten, bleibt dann immer noch gültig; und diese konstatierbare Tatsache läßt sich nur erklären, wenn die Wahrscheinlichkeitsgrenze nicht das Aushilfsmittel einer unvollkommenen Erkenntnis, sondern eine Eigenschaft des Naturgeschehens darstellen. Soll also der Wahrscheinlichkeitsbegriff eliminiert werden, so müßten die Wahrscheinlichkeitsgesetze als Folge der kausalen Gesetze erwiesen werden; aber ein solcher Beweis dürfte sicherlich unmöglich sein.

Es gelingt deshalb der Determinationshypothese keineswegs, den Wahrscheinlichkeitsbegriff entbehrlich zu machen; und darum spricht nichts dagegen, den umgekehrten Weg zu gehen, auf den Determinismus zu verzichten und den Wahrscheinlichkeitsbegriff als Grundbegriff der Erkenntnis aufzustellen. Schließlich will man doch mit der strengen Kausalhypothese nur den Gedanken ausdrücken, daß es für die Abweichungen von der strengen Gesetzmäßigkeit wieder eine kausale Erklärung geben muß; aber gerade diesen Gedanken kann man auch ohne die Hypothese einer Grenze beibehalten. Wir lassen also die Implikationshypothese gelten, und zwar nicht nur in der Form „wenn A ist, dann folgt B “, sondern auch in der umgekehrten Form „wenn B ist, so ist A vorausgegangen.“ Aber wir fügen dieser Annahme noch eine Wahrscheinlichkeitsannahme hinzu, welche sich auf die in A und B nicht mitberücksichtigten Faktoren bezieht und besagt, daß diese nach den Regeln der Wahrscheinlichkeitsrechnung zum Ausdruck kommen. Beide Annahmen sollen auf jeder Genauigkeitsstufe gelten, und wir verzichten auf die Behauptung, daß die zweite Annahme schließlich überflüssig wird. An Stelle der einheitlichen Hypothese des Determinismus begnügen wir uns also mit dem Nebeneinander zweier Annahmen: der Annahme einer kausalen Verknüpfung für die beherrschenden Faktoren des Geschehens und der Annahme einer wahrscheinlichkeitsgemäßen Verteilung für den Einfluß des Restes. Es entspricht sicherlich den Forderungen einer möglichst getreuen Naturbeschreibung, diese Doppelheit einer einheitlichen Annahme vorzuziehen, welche so wenig gerechtfertigt werden kann wie der Determinismus.

Jedoch ist es nicht einmal notwendig, die beiden Annahmen getrennt nebeneinander zu stellen. Die Aufspaltung des Geschehens in einen kausalen Teil und einen Wahrscheinlichkeitsteil ist lediglich von formaler Bedeutung; sie läßt sich ersetzen durch die eine Annahme, daß zwischen Ursache und Wirkung ein wahrscheinlichkeitsgemäßer Zusammenhang besteht. Es kann gleichgültig sein, ob A mit Gewißheit B bewirken würde, wenn weiter keine Faktoren da wären; da dieser Fall nie vorkommt, so begnügen wir uns mit der einen Annahme: „wenn A ist, so bestimmt es nach den Gesetzen der Wahrscheinlichkeit ein B .“ Der Grad der Wahrscheinlichkeit kann durch möglichst genaue Festlegung der beteiligten Faktoren beliebig nahe an 1 gesteigert werden¹⁾ — hierin drückt sich jetzt der Gedanke aus, daß zu jeder Abweichung in der Wirkung wieder eine Ursache gefunden werden kann — aber immer behält, für jede erreichbare Stufe, die Beziehung zwischen Ursache und Wirkung den Charakter einer Wahrscheinlichkeit. Wir denken uns also eine Welt, in der alle Abhängigkeiten von derselben Art sind, wie das Auftreffen einer Würfelseite mit dem Wurf im Zusammenhang steht; jeder Schritt des Geschehens ist ein Würfelspiel, und nur die große Wahrscheinlichkeit einzelner Reihen hat uns verführt, in ihnen eine sichere Gesetzmäßigkeit verborgen zu sehen. Mit dieser Auffassung sind wir dann ebenfalls zu einer einheitlichen Annahme über den Charakter des Geschehens gekommen, nur daß wir nicht die Wahrscheinlichkeitsannahme, sondern die Kausalannahme fortgelassen haben. Eine solche Welt besitzt in jedem ihrer Elemente allein einen Wahrscheinlichkeitszusammenhang.

Es ist die Forderung nach einem Minimum von Voraussetzungen, die uns zu dem Verzicht auf die strenge Kausalität

¹⁾ Es läßt sich in Zweifel ziehen, ob die Wahrscheinlichkeit in jedem Falle tatsächlich beliebig nahe an 1 gesteigert werden kann, oder ob nicht an gewissen Stellen vorher Grenzen auftreten. Diese Grenzen könnten auch praktisch unerreichbar bleiben, so daß der Satz in Geltung bliebe, daß zu jeder erreichten Genauigkeitsstufe eine höhere existiert. So berechtigt eine derartige Vermutung erscheinen mag — sie würde bestätigt werden, wenn die Quantentheorie den Versuch einer kausalen Erklärung aufgibt und sich mit den Wahrscheinlichkeitssprüngen der Elektronen begnügt — sie soll hier nicht erörtert werden, und alles folgende ist auch mit der nach 1 steigerungsfähigen Wahrscheinlichkeit verträglich.

zwingt. Jedoch werden wir finden, daß uns mit der Entwicklung der Theorie des Wahrscheinlichkeitszusammenhangs zugleich ein Erfolg zuteil wird, den sie vor dem Determinismus voraus hat und der sie darum in hohem Maße rechtfertigt: das ist die Aufklärung der Begriffe Vergangenheit und Zukunft.

Daß die Zeitordnung auf gewisse Eigenschaften der Kausalstruktur begründet werden kann, ist durch die Untersuchungen von K. Lewin,¹⁾ R. Carnap²⁾ und dem Verfasser³⁾ neuerdings klar gestellt worden. Was „früher“ und „später“ heißt, läßt sich durch Kausalreihen definieren; nur weil Ereignisse durch Kausalreihen verbunden werden können, besitzen sie ein Zeitverhältnis. Die für diese Ordnung notwendigen Eigenschaften der Kausalreihen lassen sich als Axiome formulieren; unter ihnen spielt der Ausschluß geschlossener Kausalreihen eines einzigen Richtungsinns eine wichtige Rolle. So läßt sich eine Topologie der Zeit gewinnen, in der die Grundbegriffe „früher“, „später“, „gleichzeitig“ definiert werden. Aber was damit bisher nicht gelöst werden konnte, ist das Problem des „jetzt“.

Was heißt „jetzt“? Plato lebte früher als ich, und Napoleon VII. wird später leben als ich. Aber wer von diesen dreien lebt jetzt? Zweifellos habe ich ein deutliches Gefühl dafür, daß ich jetzt lebe. Aber hat diese Aussage einen objektiven Sinn über mein subjektives Erlebnis hinaus? Ihre Bedeutung könnte sich in der Schilderung eines psychologischen Zustandes erschöpfen. Aber ist es nicht doch möglich, ihr eine objektive Bedeutung zu geben?

Man wird zunächst versuchen, diese objektive Bedeutung in einer Aussage über Gleichzeitigkeitsbeziehungen zu finden. Dann ist die Aussage „ich lebe jetzt“ identisch mit Aussagen der Form „ich lebe gleichzeitig mit Herrn A“ oder „ich lebe gleichzeitig mit dem und dem Ereignis.“ Ist dies der Fall, so gibt es kein besonderes „jetzt“, sondern die Bedeutung dieses Wortes ist zu-

¹⁾ Zschr. f. Phys. 13, 62, 1923.

²⁾ Kantstudien 1925, 30, S. 331.

³⁾ Axiomatik der relativistischen Raum-Zeit-Lehre, Vieweg 1924, und Physikal. Zschr. 1921, 22, 683. Dieser Axiomatik nahe verwandt ist die „Axiomatik der speziellen Relativitätstheorie“ von C. Carathéodory, Berl. Akad. Ber. 1924, S. 12.

rückführbar auf die Begriffe „früher“, „später“, „gleichzeitig“. Aber erschöpft sich damit der Sinn des „jetzt“?

Wenn das Jetzt auf Gleichzeitigkeit zurückführbar ist, so würde der Sinn der Frage „was geschieht jetzt?“ in folgendem bestehen: diese Frage stellt selbst ein Ereignis F vor, das im Weltablauf seine Position hat, und gefragt wird nach dem, was mit F gleichzeitig ist. Dennoch ist diese Antwort, obzwar richtig, nicht erschöpfend. Denn es steht nicht in meiner Hand, die Position dieses F auszusuchen; dieses F ordnet sich von selbst in den Zeitpunkt „jetzt“ ein. Wenn man antwortet, daß ich die Lokalisation des F sehr wohl in der Hand habe, indem ich mit der Stellung der Frage warten kann, so ist hierauf folgendes zu erwidern. Ich kann jedenfalls nicht alle Zeitpunkte dafür aussuchen, sondern nur zukünftige. Der Zeitpunkt aber, welcher diese wählbaren Zeitpunkte von den nicht wählbaren trennt, ist das Jetzt. Man kann eben mit solchen Versuchen nicht dem Zwang entrinnen, der für uns einen Jetzt-Punkt als Erlebnis der Grenze zwischen Vergangenheit und Zukunft absolut auszeichnet.

Das Problem läßt sich deshalb auch formulieren als die Frage nach dem Unterschied von Vergangenheit und Zukunft. Für den Determinismus gibt es einen solchen Unterschied nicht. Wenn in irgend einem zeitlichen Querschnitt die Zukunft bereits völlig bestimmt ist, so macht es keinen Unterschied, ob sie schon abgelaufen ist oder noch ablaufen wird. Der Ablauf bringt nichts Neues; das was in 100 Jahren geschehen wird, ist mir in demselben Sinne gegeben wie die Ereignisse des vergangenen Krieges, und ich könnte mich in grundsätzlich derselben betrachtenden Weise über die Kriege Napoleons VII. unterhalten wie über die Kämpfe bei Verdun. Dann besteht in bezug auf das „jetzt“ kein Unterschied zwischen Plato und mir; ich kann ebensogut sagen, Plato lebt jetzt, und ich bin noch Zukunft. Zwar, daß Plato früher lebt als ich, könnte ich dann aussagen, denn ein „früher“ und „später“ gibt es auch für den Determinismus. Aber es gibt kein „jetzt“; es gibt keinen ausgezeichneten Zeitpunkt, und das Gefühl, daß mein Dasein eine Realität ist, Platos Leben aber nur noch seine Schatten in die Realität wirft, muß ein Irrtum sein. Dem widerspricht aber die ganze Haltung unseres Daseins, wir haben eine vollkommen verschiedene Einstellung gegen die

Zukunft als gegen die Vergangenheit; und wenn man nicht jede einzelne unserer Handlungen, jeden Gedanken, der uns in der Gestaltung unseres täglichen Lebens begleitet, als einen einzigen großen Irrtum auffassen will, muß der Determinismus falsch sein.

Es soll damit nicht gesagt sein, daß der Determinismus falsch ist; aber über den Gegensatz muß man sich ganz klar sein. Hat der Determinismus recht, so ist es durch nichts zu rechtfertigen, daß wir uns für den morgigen Tag eine Handlung vornehmen, für den gestrigen Tag aber nicht. Es ist wohl wahr, daß wir dann gar nicht die Möglichkeit haben, auch nur den Vorsatz zu der morgigen Handlung und den Glauben an Freiheit zu unterlassen — gewiß nicht, aber einen Sinn hat unser Tun dann nicht. Denn dann ist der morgige Tag heute schon in demselben Sinne vorbei wie der gestrige. Aber zu dieser Konsequenz zwingt eben nur der Determinismus — wenn man auf ihn verzichtet, läßt sich der Widerspruch zu unserm elementaren Lebensgefühl vermeiden. Gewiß darf ein solches Gefühl nicht entscheiden, wenn der Verstand überzeugend dagegen spricht — aber man analysiere zuvor den Verstand, ob seine Behauptung notwendig ist. Und das ist sie nicht.

Denn entschließt man sich zur Theorie des Wahrscheinlichkeitszusammenhangs, so ergibt sich gerade der Unterschied zwischen Vergangenheit und Zukunft, der unserm Gefühl entspricht. Gibt es keine völlige Bestimmtheit des Geschehens, so kann man nicht sagen, daß die Zukunft jetzt schon feststeht. Das Gegenteil des Berechneten ist dann auch immer möglich. Die Vergangenheit dagegen steht fest, und die Gegenwart ist diejenige Schwelle, auf welcher die Welt vom Zustand der Unbestimmtheit in den der Bestimmtheit übergeht. Es ist also im Zustand der Welt ein Querschnitt ausgezeichnet, den man Gegenwart nennt; das „jetzt“ hat eine objektive Bedeutung. Auch wenn kein Mensch mehr lebt, gibt es ein „jetzt“; der „jetzige Zustand des Planetensystems“ ist dann eine ebenso bestimmte Angabe wie „der Zustand des Planetensystems zur Zeit von Christi Geburt.“

In dem vierdimensionalen Bild der Welt, wie es etwa die Relativitätstheorie benutzt, gibt es einen solchen ausgezeichneten Querschnitt nicht. Aber das liegt nur daran, daß in diesem Bild ein wesentlicher Inhalt weggelassen ist. Sollen irgend welche

Aussagen über die Welt, über vergangene oder zukünftige Ereignisse, gemacht werden, so müssen sie an gewisse wahrgenommene Ereignisse durch Schlußketten angeknüpft werden. Und zwar müssen alle diese Anknüpfungspunkte auf einem Querschnitt liegen, eben dem Gegenwartsquerschnitt. In der Tat: will ich wissen, wann Karl der Große geboren wurde, so muß ich ein Geschichtswerk aufschlagen; die Wahrnehmung der Zahl ist das Gegenwartereignis, von dem erst eine Schlußkette zu der Behauptung führt, daß sie das Geburtsjahr Karls des Großen bedeutet. (Die Schlußkette enthält z. B. die Annahme, daß das Buch ein hinreichend zuverlässiges Geschichtswerk ist.) Will ich eine Sonnenfinsternis berechnen, so müssen entweder wieder gedruckte Zahlen eines Buches, die ich „jetzt“ lese, oder gegenwärtige Beobachtungen von Sonne und Mond als Ausgang dienen. Zahlen, die ich nicht nachlese, sondern in der Erinnerung habe, müssen „jetzt“ gewußt werden; das Erinnerungserlebnis ist hier der Wahrnehmung vergleichbar und führt auch nur durch Schlußketten (z. B. Kontrollen der Sicherheit des Gedächtnisbildes) zur behaupteten Tatsache. Es gibt also zu einem gegebenen Weltzustand einen Querschnitt derart, daß alle Aussagen an ihn angeknüpft werden müssen, sowohl über die Vergangenheit als über die Zukunft.

Obgleich wir diesen Querschnitt dadurch charakterisiert haben, daß wir alle Aussagen über die Welt an ihn anknüpfen müssen, wird er dadurch nicht subjektiv definiert. Denn es liegt nicht an uns, daß wir ihn wählen müssen, sondern gerade an dem Zustand der Welt. Zu jedem gegebenen Weltzustand gibt es einen ausgezeichneten Querschnitt. Das Weltbild der Relativitätstheorie sollte richtig wie in Fig. 1a gezeichnet werden, und der Ablauf der Welt besteht darin, daß der Zustand der Fig. 1a in den der Fig. 1b usw. übergeht.¹⁾ Man kann den Weltablauf nicht in einem Bild zeichnen, sondern nur in einer Folge derartiger Bilder wie Fig. 1. Es ist nur eine (für viele Zwecke natürlich zulässige) Vereinfachung, wenn man überall den Querschnitt und die Pfeilspitzen wegläßt und die Folge durch ein einziges Bild ersetzt.

¹⁾ Dabei ist in Fig. 1b derjenige Teil, der in Fig. 1a der Zukunft entspricht, etwas verändert gezeichnet; es sei damit angedeutet, daß die Zukunft anders eingetroffen ist, als sie in 1a berechnet wurde.

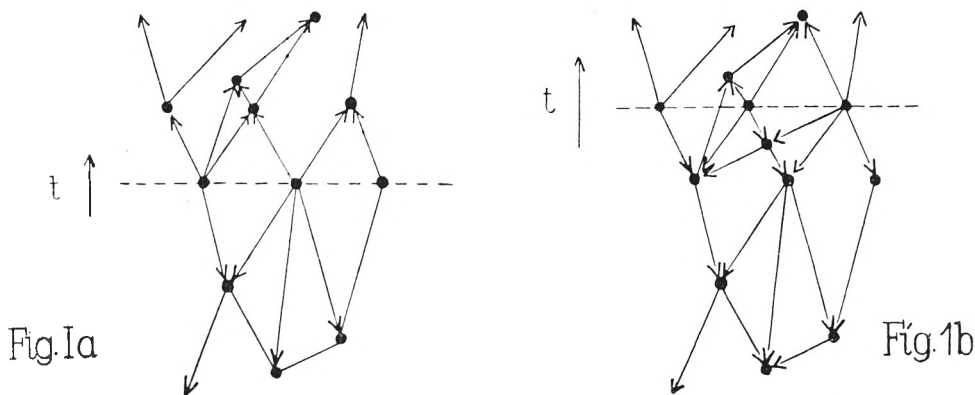


Fig. 1. Der Weltablauf als Folge von Strukturbildern mit ausgezeichnetem Gegenwartsquerschnitt.

Obgleich wir von einem ausgezeichneten Querschnitt sprechen, soll damit nicht etwa die Existenz einer absoluten Gleichzeitigkeit behauptet werden. Sondern wir müssen unsere Aussage im Sinne der Relativitätstheorie korrigieren: die Richtung des ausgezeichneten Querschnitts ist innerhalb eines gewissen Intervalls willkürlich. Wir können dies auch dann noch zulassen, wenn wir das Jetzt durch das subjektive Erlebnis „jetzt“ definieren. Die Erlebnisse eines einzelnen Menschen stellen im Kausalschema nur einen Querschnitt von sehr geringer Breite dar, den man als nahezu punktförmig betrachten kann. Dann ist jeder im Sinne der Relativitätstheorie zulässige Gleichzeitigkeitsschnitt durch dieses Punktereignis ein zulässiger Jetzt-Schnitt. Der Jetzt-Schnitt läßt sich also von einem Punkt aus definieren. Das stimmt auch überein mit der in Fig. 1 gegebenen Definition des Jetzt-Schnitts durch die Umkehr der Pfeilrichtung. Der von einem Punktereignis P nach vorwärts und rückwärts ausgehende Wirkungskegel $ds^2 = 0$ zerteilt die Welt bereits derart, daß alle in P zusammenlaufenden Wirkungslinien mit Pfeilspitzen im Sinne unserer Fig. 1 versehen werden können. Ungeordnet bleiben dabei nur die Punkte des Zwischengebiets (im Minkowskischen Sinne), und dieses Gebiet wird eben durch die zulässigen Jetzt-Schnitte durch P ausgefüllt. Wenn wir im folgenden von dem ausgezeichneten Querschnitt reden, meinen wir genauer einen beliebigen der ausgezeichneten Schnitte, und auch die Fig. 1 ist in diesem Sinne zu verstehen.

In der Art, wie sich Vergangenheit und Zukunft von dem ausgezeichneten Querschnitt aus bestimmen, unterscheiden sich beide. Wir wollen dies jetzt mit Hilfe der Theorie des Wahrscheinlichkeitszusammenhangs zeigen und dabei klar legen, in welchem Sinne die Vergangenheit „objektiv bestimmt“, die Zukunft „objektiv unbestimmt“ genannt werden kann. Wir lassen jedoch dabei die Vorstellung, daß die Welt mit einem stetigen Feld erfüllt ist, fallen, und denken uns einzelne Ereignisse (die Knotenpunkte in Fig. 1), die durch Schlußketten miteinander verknüpft sind. Diese Auffassung ermöglicht es uns, den Weltzusammenhang auf die topologischen Eigenschaften einer Netzstruktur zu begründen. Die Ausdehnung der Theorie auf stetige Felder ist mit Schwierigkeiten verknüpft, die sich vorläufig noch nicht beseitigen lassen.

II. Topologie der Wahrscheinlichkeitsimplikation.

Die Relation, welche an Stelle der strengen kausalen Verknüpfung der Ereignisse tritt, nennen wir Wahrscheinlichkeitsimplikation. Liegt ein Ereignis A vor, so beobachten wir, daß dann mit einer gewissen Regelmäßigkeit auch das Ereignis B auftritt. Es braucht nicht immer B aufzutreten, aber die Fälle des Auftretens und Nichtauftretens von B sind nach den Gesetzen geregelt, die in der Wahrscheinlichkeitsrechnung niedergelegt sind. Dabei gehört zu diesen Gesetzen nicht nur die Regelmäßigkeit des Häufigkeitsverhältnisses von Eintreffen und Nichteintreffen, sondern auch die Regelmäßigkeit in den Abweichungen von diesem Verhältnis, d. h. die Gesetze der Streuung. Wir sagen dann

$$A \rightarrow B$$

gesprochen: „ A impliziert mit Wahrscheinlichkeit B “, oder auch: „ A bestimmt B “. Dabei soll über den Grad der Wahrscheinlichkeit nichts gesagt sein, dieser kann zwischen 0 und 1 (einschließlich) liegen; die Relation $A \rightarrow B$ gilt also nicht nur dann, wenn nach dem Sprachgebrauch B „wahrscheinlich gemacht“ wird durch A , sondern auch, wenn B „unwahrscheinlich gemacht“ wird durch A . Das Zeichen $\rightarrow$ für die Wahrscheinlichkeitsimplikation geht aus dem Zeichen $\supset$ der strengen (logischen) Im-

plikation hervor, indem der Querstrich hinzutritt. Die strenge Implikation geht als Grenzfall aus der Wahrscheinlichkeitsimplikation hervor, wenn die Wahrscheinlichkeit $= 1$ wird.

Man wird gegen die Einführung der Wahrscheinlichkeitsimplikation zwei Einwände machen. Erstens: Wie ist es möglich, die Regelmäßigkeit des Häufigkeitsverhältnisses von A und B für alle Fälle zu behaupten, wenn sie doch nur für eine endliche Zahl von Fällen beobachtet ist? Wir antworten, daß wir auf diese Frage, die das Problem der Induktion darstellt, hier nicht eingehen wollen, sondern daß wir es als möglich und sinnvoll voraussetzen wollen, von einer endlichen Zahl von Beobachtungen auf alle Beobachtungen mit Wahrscheinlichkeit zu schließen. Diese Voraussetzung macht nicht nur unsere Wahrscheinlichkeitstheorie, sondern jede wissenschaftliche Naturerkenntnis; wir wollen ihre Berechtigung deshalb unterstellen. Zweitens wird man einwenden: Was für einen Sinn hat es, dem Eintreffen des einzelnen Ereignisses B eine Wahrscheinlichkeit zuzuschreiben, wenn diese Zahl doch gar nichts für den Einzelfall, sondern nur etwas für beliebig lange Wiederholungsreihen bedeutet? Hierauf antworten wir ebenfalls, daß wir diese Aussage als sinnvoll voraussetzen wollen; und auch diese Voraussetzung gilt nicht nur für unsere Theorie, sondern wird in der Wissenschaft und im täglichen Leben ständig gemacht. Ihre Kritik — welche zu beachten hat, daß das Problem für den Einzelfall grundsätzlich nicht anders liegt als für jede endliche Anzahl vorauszusagender Fälle — ist eine sehr wichtige Frage der Erkenntnistheorie, aber sie soll uns hier nicht beschäftigen.

Wir betrachten hier also die Wahrscheinlichkeitsimplikation als einen Grundbegriff, ähnlich wie man in der Logik die Implikation als nicht ableitbaren Grundbegriff einführen kann. $A \Rightarrow B$ bedeutet: „Wenn A ist, so ist mit Wahrscheinlichkeit B .“ Oder auch: „Wenn mit Wahrscheinlichkeit A ist, so ist mit Wahrscheinlichkeit B .“ Aber was dies heißt: „ B ist mit Wahrscheinlichkeit“ nehmen wir als nicht weiter zerlegbaren Grundbegriff an. Nicht zwischen irgend welchen Ereignissen, sondern nur zwischen gewissen Ereignissen dürfen wir die Beziehung „Wahrscheinlichkeitsimplikation“ ansetzen; welche dies sind, lehrt die Erfahrung. $A \Rightarrow B$ ist also eine materiale Aussage.

Die merkwürdigste Eigenschaft der Wahrscheinlichkeitsimplikation im Gegensatz zur Implikation besteht nun darin, daß mit $A \rightarrow B$ stets auch $A \rightarrow \bar{B}$ gegeben ist, wo $\bar{B}$ (gesprochen: non- B) das Fehlen des Ereignisses B bedeutet. Das ist vom Standpunkt der Wahrscheinlichkeitsrechnung selbstverständlich; ist p das Maß der Wahrscheinlichkeit, mit der B bestimmt wird, so ist $1 - p$ das entsprechende Maß für $\bar{B}$. Mit der Regelmäßigkeit des Eintreffens von B ist auch die entsprechende Regelmäßigkeit für das Fehlen von B , also das Eintreffen von $\bar{B}$, gegeben. Wir können diese Grundeigenschaft so schreiben

$$(A \rightarrow B) \supset (A \rightarrow \bar{B}), \quad (1)$$

wo $\supset$ die strenge Implikation bedeutet.

Von der Behauptung $A \rightarrow \bar{B}$ ist die Behauptung

$$\overline{A \rightarrow B}$$

wohl zu unterscheiden. Diese besagt, daß es falsch ist, wenn man $A \rightarrow B$ behauptet. Dies bedeutet, daß keine Regelmäßigkeit zwischen dem Stattfinden von A und B besteht, wie sie die Wahrscheinlichkeitsgesetze verlangen. Mit Satz (1) ergibt sich sogleich

$$\overline{(A \rightarrow B)} \supset (A \rightarrow \bar{B}).$$

Es gibt gewisse Fälle, in denen außer $A \rightarrow B$ auch $B \rightarrow A$ gilt. Hier ist also die Wahrscheinlichkeitsimplikation umkehrbar. Ob Umkehrbarkeit vorliegt, kann nur die Erfahrung lehren; es ist also wieder eine materiale Behauptung. Das Maß der Wahrscheinlichkeit ist im allgemeinen für beide Richtungen verschieden.

Einige Beispiele: Das Steigen des Barometers impliziert mit Wahrscheinlichkeit, daß das Wetter gut wird. Umgekehrt: Wenn das Wetter gut wird, ist mit Wahrscheinlichkeit zu folgern, daß das Barometer gestiegen ist. Dagegen: Wenn ich Herrn X auf der Straße Y treffe, so folgt mit Wahrscheinlichkeit, daß Herr X nach Z geht. Das Umgekehrte gilt nicht: „Wenn Herr X nach Z geht, so folgt nicht, auch nicht mit Wahrscheinlichkeit, daß ich Herrn X auf der Straße Y treffe.“

Im folgenden wird eine Zusammenstellung von Gesetzen der Wahrscheinlichkeitsimplikation gegeben, die weder den Anspruch macht, vollständig zu sein, noch den, eine Tafel unabhängiger Axiome zu bedeuten. Doch dürften die wichtigsten Gesetze da-

mit getroffen sein. Wir bedienen uns dabei der Russelschen Schreibweise der mathematischen Logik, nur mit dem Unterschied, daß wir an Stelle des Russelschen Zeichens für die Negation die übersichtlichere Überstreichung wählen.

Es bedeutet also:

$a \supset b$ a impliziert b

$a \rightarrow b$ a impliziert mit Wahrscheinlichkeit b

a, b a und b

$a \vee b$ a oder b oder beides (das nicht ausschließende „oder“)

$\bar{a}$ non- a (Verneinung).

Gesetze der Wahrscheinlichkeitsimplikation.

- | | |
|---|---|
| 1*. $(a \rightarrow b) \supset (a \rightarrow \bar{b})$ | Doppeldeutigkeit |
| 2*. $(a \rightarrow b, c) \supset (a \rightarrow b) \cdot (a \rightarrow c)$ | Auflösung des hinteren „und“ |
| 3*. $(a \vee b \rightarrow c) \supset (a \rightarrow c) \cdot (b \rightarrow c) \cdot (a, b \rightarrow c)$ | Auflösung des vorderen „oder“ |
| 4*. $(a \rightarrow b \vee c) \supset (a \rightarrow b) \vee (a \rightarrow c)$ | Auflösung des hinteren „oder“ |
| 5*. $(a \rightarrow b) \supset (a, c \rightarrow b)$ | Faktor vorn |
| 6*. $(a \rightarrow b) \cdot (a \rightarrow c) \supset (a \rightarrow b, c)$ | hintere Multiplikation |
| 7*. $(a \rightarrow c) \cdot (b \rightarrow c) \supset (a \vee b \rightarrow c)$ | vordere Addition |
| 8*. $(a \rightarrow b) \vee (a \rightarrow c) \supset (a \rightarrow b \vee c)$ | hintere Addition |
| 9*. $(a \rightarrow b) \cdot (b \rightarrow c) \supset (a \rightarrow c)$ | Transitivität |
| 10*. $(a \rightarrow b) \cdot (b \rightarrow c) \supset (a \rightarrow c)$ | Transitivität für den partiellen Grenzfall. |

Die Gesetze der Wahrscheinlichkeitsimplikation sind denen der Implikation ganz analog. Man muß dabei nur beachten, daß in ihnen neben der Wahrscheinlichkeitsimplikation auch die Implikation auftritt; man kann nicht etwa in einem für die Implikation richtigen Satz überall das Zeichen $\supset$ durch $\rightarrow$ ersetzen, sondern darf dies nur an gewissen Stellen tun. Die Wahrscheinlichkeitsimplikation ist also auf die Implikation basiert, und die strenge Implikation kann durch die Wahrscheinlichkeitsimplikation nicht entbehrlich gemacht werden. Umgekehrt darf man auch nicht verlangen, obgleich die Implikation ein Grenzfall der Wahrscheinlichkeitsimplikation ist, daß alle Gesetze richtig bleiben,

wenn man überall $\rightarrow$ durch $\supset$ ersetzt. Auch diese Einsetzung darf nur an gewissen Stellen geschehen. Setzt man z. B. in (1*) für das erste $\rightarrow$ ein $\supset$ ein, so darf man dies für das andere $\rightarrow$ nicht tun. Denn das Maß dieser zweiten Wahrscheinlichkeitsimplikation wird $= 0$, wenn das der ersten $= 1$ wird. — Die Bedeutung einzelner dieser Gesetze wird erst im folgenden bei den Anwendungen klar werden.

Wir werden nun daran gehen, mit Hilfe der Wahrscheinlichkeitsimplikation Aussagen über die Kausalstruktur der Welt zu machen. Es sind dies topologische Aussagen, weil in der Wahrscheinlichkeitsimplikation und ihren bisher gegebenen Gesetzen noch kein Gebrauch von dem Maß der Wahrscheinlichkeit gemacht wird. Wir denken uns auf dem Gegenwartsquerschnitt gewisse Ereignisse gegeben, und schließen aus ihnen mit Hilfe von Naturgesetzen auf andere Ereignisse. Die Naturgesetze haben alle die Form $a \rightarrow b$. Woher wir diese Gesetze im einzelnen kennen, insbesondere wie es möglich ist, derartige Gesetze zwischen zeitlich folgenden Ereignissen zu finden, wenn doch immer nur gleichzeitige Ereignisse gegeben sind — das soll uns hier zunächst nicht interessieren. Wir nehmen die Gesetze also als gegeben an. Wir wollen zeigen, daß die Schlußweise topologisch eine andere ist, je nachdem auf vergangene oder zukünftige Ereignisse geschlossen wird.

Für die Aufdeckung des strukturellen Unterschieds der Zeitrichtung benutzen wir folgendes Verfahren. Wir nehmen zunächst an, daß uns anderweitig bekannt sei, ob auf vergangene oder zukünftige Ereignisse geschlossen wird; so gewinnen wir die Charakterisierung der Schlußweise. Umgekehrt darf dann die Besonderheit der Schlußweise zur Definition der Zeitrichtung benutzt werden.

Die einfachste Ordnung der Ereignisse ist die ungeteilte Kette.



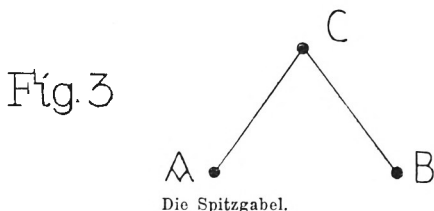
Die ungeteilte Kette.

Hier gilt

$A \rightarrow B$	$B \rightarrow C$	$A \rightarrow C$	(2)
$C \rightarrow B$	$B \rightarrow A$	$C \rightarrow A$	

Hier ist keine Richtung ausgezeichnet. Die ungeteilte Kette liefert also keine Kennzeichnung der Zeitrichtung; dies gelingt erst mit dem Auftreten von Knotenpunkten. Wir werden deshalb dazu geführt, die Zeitordnung auf die Eigenschaften einer Netzstruktur zu begründen.

Die einfachste Grundform einer Netzstruktur nach Fig. 1 ist die Gabel. Wir betrachten zunächst eine mit der Spitze in die Zukunft weisende Gabel, die wir Spitzgabel nennen; zum Unterschied von der später zu besprechenden disjunktiven Spitzgabel heißt sie auch konjunktive Spitzgabel.



$$\begin{array}{lll} \text{Für sie gilt } A.B \rightarrow C & C \rightarrow A.B & \\ \hline A \rightarrow C & C \rightarrow A & A \rightarrow B \\ B \rightarrow C & C \rightarrow B & B \rightarrow A \end{array} \quad (3)$$

Das Charakteristische ist hier das vordere „und“ in der ersten links stehenden Aussage. Es ist, wie bei der strengen Implikation, nicht auflösbar; d. h. nur A und B zusammen bestimmen C . Dies ist das Charakteristische eines Schlusses in die Zukunft. Es gilt also: $A \rightarrow C$.

Beispiel: In A und B wird je eine Billardkugel losgeschleudert, C ist das Ereignis ihres Zusammenstoßes. Eine Wahrscheinlichkeit für C ist erst gegeben, wenn beide Ereignisse A und B stattfinden, und kann nur aus beiden Einzelwahrscheinlichkeiten für das Eintreffen der Kugeln an dem Ort C berechnet werden. Ist über den Abgang der Kugel in B nichts bekannt, d. h. existiert keine Wahrscheinlichkeit für das Eintreffen dieser Kugel an dem Ort von C , so besteht zwischen dem Abgang der Kugel in A und dem Ereignis C kein Wahrscheinlichkeitszusammenhang.

Die Aussage $A.B \rightarrow C$ ist umkehrbar; dadurch entsteht $C \rightarrow A.B$. Dieses hintere „und“ ist nach 2* auflösbar; so entstehen $C \rightarrow A$ und $C \rightarrow B$. Wir können hier bereits die Regel für die Zeitrichtung gewinnen:

Richtungsregel: Wenn die Wahrscheinlichkeitsimplikation nur in einer Richtung gilt, so ist das vordringende Ereignis das zeitlich spätere.

In Zeichen:

$$(C \rightarrow A) \cdot (A \rightarrow C) \supset (A < C), \quad (4)$$

wo $A < C$ „ A ist früher als C “ bedeutet.

Da die Wahrscheinlichkeitsimplikation zwischen C und A nur in einer Richtung gilt, läßt sich eine Implikation zwischen A und B in keiner Richtung herstellen; denn $A \rightarrow B$ würde nach 9* $A \rightarrow C$ und $C \rightarrow B$ verlangen, und entsprechend würde $B \rightarrow A$ voraussetzen $B \rightarrow C$ und $C \rightarrow A$. Die Spitzgabel ist also intransitiv.

Durch die Aussagen (3) ist die Spitzgabel völlig festgelegt. Seien irgend drei Ereignisse gegeben und sei es bekannt, daß zwischen ihnen die Beziehungen (3) gelten, so müssen diese Ereignisse eine Spitzgabel bilden. Welches der drei Ereignisse die in die Zukunft weisende Spitze darstellt, ist an dem unsymmetrischen Auftreten dieses Ereignisses in den Beziehungen (3) kenntlich. Vertauscht man in (3) A mit B , so entsteht dasselbe System von Sätzen wie vorher. Vertauscht man aber C mit A oder B , so entstehen andere Sätze. Die intransitive Gabel hat also eine topologisch ausgezeichnete Ecke.

Die mit der Spitze in die Vergangenheit weisende Gabel soll Sattलगabel genannt werden; auch hier werden wir später eine disjunktive Sattलगabel von der jetzt zu besprechenden konjunktiven Sattलगabel unterscheiden. Für sie gelten die Relationen:

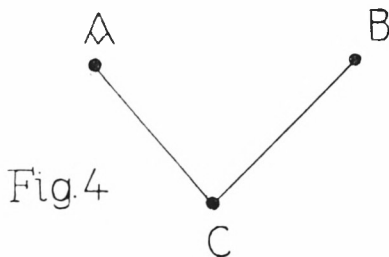


Fig. 4

Die Sattलगabel.

$$\begin{array}{lll}
 A \vee B \rightarrow C & C \rightarrow A \cdot B & \\
 A \rightarrow C & C \rightarrow A & A \rightarrow B \\
 B \rightarrow C & C \rightarrow B & B \rightarrow A
 \end{array} \quad (5)$$

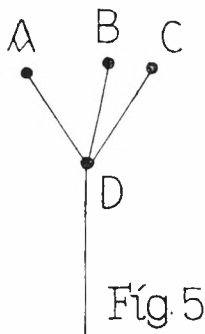
Das Charakteristische ist hier das vordere „oder“ in der ersten links stehenden Aussage. Es ist nach 3* auflösbar und führt darum, im Gegensatz zur Spitzgabel, auf $A \rightarrow C$ und $B \rightarrow C$. Mit 9* gilt infolgedessen auch $A \rightarrow B$, und auch $B \rightarrow A$; die Sattलगabel ist also transitiv.

Beispiel: A und B mögen wieder das Abschleudern je einer Billardkugel bedeuten; aber C bedeutet hier die gemeinsame Ursache, etwa das Signal, auf welches hin die beiden Kugeln losgeschleudert werden. Beobachte ich nur A , so darf ich bereits mit Wahrscheinlichkeit schließen, daß das Signal gegeben wurde. Auch wenn B gar nicht stattfindet, darf von A mit Wahrscheinlichkeit auf C geschlossen werden; es hat vielleicht der Mechanismus des Abschleuderns in B versagt. Die zu der Beobachtung von A hinzu kommende Beobachtung von B verstärkt nur die Wahrscheinlichkeit für C . Und beobachte ich A , während mir über das Stattfinden von B nichts bekannt ist, so darf ich von A über die gemeinsame Ursache C auf B schließen.

Wir erkennen hier den entscheidenden Unterschied zwischen Vergangenheit und Zukunft. Die Spitzgabel und die Sattलगabel sind symmetrisch in Bezug auf den Gegenwartsquerschnitt; wäre die Schlußweise in die Vergangenheit dieselbe wie in die Zukunft, so müßten die Relationen (3) und (5) identisch sein. Aber sie sind gerade in einem wesentlichen Punkt unterschieden: Der Schluß in die Zukunft verlangt ein vorderes „und“, der Schluß in die Vergangenheit braucht nur ein vorderes „oder“. Den Schluß in die Zukunft erlaubt nur die Gesamtheit aller Ursachen, aber in die Vergangenheit kann man schon aus einer Teilwirkung schließen.

Dabei erfolgt die Kennzeichnung der Zeitrichtung gerade durch die intransitive Gabel, wie wir es in der Richtungsregel formuliert haben; die transitive Gabel ermöglicht die Kennzeichnung der Richtung nicht. Denn gerade wegen der Transitivität hat diese Gabel keine topologisch ausgezeichnete Ecke. Vertauscht man in (5) A mit B , oder B mit C , oder C mit A , so entstehen wieder die Sätze (5) oder solche, die aus ihnen nach den aufgeführten Gesetzen der Wahrscheinlichkeitsimplikation hervorgehen. Darum kann aus (5) nicht gefolgert werden, welche Ecke die zeitlich zurückliegende ist. Es kann überhaupt nicht ge-

folgert werden, daß die eine Ecke in der Vergangenheit liegen muß. Denkt man sich von einem Ereignis D drei in die Zukunft weisende Kausalketten ausgehend (Fig. 5), die zu den Ereignissen A , B , C der Gegenwart führen, so gelten zwischen diesen auch ge-



Gabel mit drei Zweigen.

rade die Beziehungen (5). Auch die ungeteilte Kette (Fig. 2) führt auf dieselben Beziehungen, denn die Beziehungen (2) sind mit (5) identisch. Darum kann aus dem Bestehen der Beziehungen (5) nicht geschlossen werden, daß eine Sattelfabel vorliegt. Über die Zeitrichtung solcher Ereignisse, die durch (5) verbunden sind, entscheidet erst ihr Zusammenhang mit Spitzgabeln im Netzwerk der Struktur.

Dagegen besteht die Bedeutung der Sattelfabel gerade in der Transitivität. Denn sie ermöglicht die Herstellung der Wahrscheinlichkeitsimplikation zwischen Ereignissen, die nicht durch eine ständig steigende oder ständig fallende Kausalkette verbunden sind. Die Herstellung der Wahrscheinlichkeitsimplikation zwischen Ereignissen desselben Gegenwartsquerschnitts gelingt deshalb nur auf dem Wege über vergangene Ereignisse, nicht über zukünftige Ereignisse; denn nur die in die Vergangenheit zeigende Gabel ist transitiv. Nur die gemeinsame Ursache, nicht die gemeinsame Wirkung stellt eine Wahrscheinlichkeitsbeziehung zwischen gleichzeitigen Ereignissen her.

Die praktische Bedeutung der Sattelfabel für die experimentelle Physik ist außerordentlich groß. Die große Mehrzahl aller Schlüsse, auch über zukünftige Ereignisse, wird auf dem Weg über die Sattelfabel gewonnen. Wird z. B. eine Temperatur im elektrischen Ofen durch die Heizstromstärke kontrolliert, so liegt eine Sattelfabel vor. Beobachtet wird ein Zeigerausschlag, von ihm wird auf die Stärke des elektrischen Stroms als Ursache rückgeschlossen, und von da wieder auf die zweite Wirkung des Stromes, die Erwärmung. Auf diesem Prinzip beruhen alle Meßinstrumente. Dabei wird die beobachtete Teilwirkung A der Ursache C so ausgewählt, daß für die Relation $A \rightarrow C$ eine sehr hohe Wahrscheinlichkeit gilt. Damit ist dann C sicher gestellt. Wird nun B beobachtet, so kann damit die Relation $C \rightarrow B$

experimentell gewonnen werden. Meistens wird auch B nicht direkt beobachtet, sondern wieder nur eine Teilwirkung D von B , welche den Schluß $D \rightarrow B$ mit hoher Wahrscheinlichkeit erlaubt. Das Schlußschema entspricht dann Fig. 6, in der die Ketten großer Wahrscheinlichkeit stark gezeichnet sind. Will man z. B. einen elektrischen Ofen auf Temperatur eichen, so bedeutet in Fig. 6

- A Zeigerausschlag am Ampèremeter
- C Stromstärke
- B Temperatur
- D Zahlenangabe eines Thermometers.

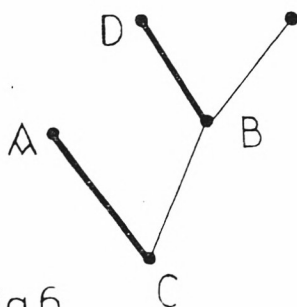


Fig. 6

Ein Analogon dieses Schlußverfahrens auf einem ganz andern Gebiet ist der Indizienbeweis der

Jurisprudenz. Bei der durch Indizien nachgewiesenen Täterschaft

an einem Mord würde in Fig. 6 etwa bedeuten:

- A Fingerabdruck des X , gefunden am Tatort
- C Anwesenheit des X am Tatort
- B Mord
- D Spuren des Mordes.

Bei diesem Schluß wird für $C \rightarrow B$ eine große Wahrscheinlichkeit angenommen, dagegen für $B \rightarrow C$ eine kleine. Aus den Spuren des Mordes allein kann man also nicht auf die Täterschaft des X schließen, wohl aber, wenn der Fingerabdruck hinzukommt.

Die Transitivität der Satteltabel gibt uns auch die Antwort auf eine Frage, die wir oben berührt haben. Gegeben ist stets nur der Gegenwartsquerschnitt, beobachtet werden also nur Wahrscheinlichkeitsimplikationen zwischen gleichzeitigen Ereignissen. Wie ist es möglich, Aussagen der Form $C \rightarrow A$ zu gewinnen, wenn C die Ursache von A ist? Auch hier ist wieder die Satteltabel das Hilfsmittel, und zwar werden vor allem Satteltabellen benutzt, deren einer Zweig eine hohe Wahrscheinlichkeit gewährt. Aber es ist in der Tat wichtig, sich darüber klar zu sein, daß

Schluß auf Ereignisse einer andern Kette mit Hilfe der Satteltabel.

die gesamte Vergangenheit eine Netzkonstruktion ist, die allein an Wahrscheinlichkeitsimplikationen zwischen gleichzeitigen Ereignissen angeknüpft wird.

Die Transitivitätseigenschaft der Sattलगabel ermöglicht es häufig, einen Schluß in die Vergangenheit mit sehr viel größerer Sicherheit auszuführen als in die Zukunft. In Fig. 7 ist eine nach Vergangenheit und Zukunft symmetrische Verkettung gezeichnet, in der die Kette AB besonders unsicher sein soll, und zwar nach beiden Richtungen. Infolge dessen ist die Vorausbestimmung des B von A aus unsicher, und ebenso die Rückbestimmung des A von B aus. Jedoch gibt es die Möglichkeit, von einem andern Ereignis C aus A als vergangenes Ereignis mit größerer Sicherheit zu bestimmen, auf dem Wege $CDEFA$. Die entsprechende Möglichkeit, B als zukünftiges Ereignis von F aus sicherer über $FEDCB$ zu bestimmen, existiert aber nicht, denn über C hinweg kann nicht geschlossen werden, weil DCB eine Spitzgabel ist. So kommt es, daß die Vorausbestimmung von Ereignissen durch eine unsichere Kette sehr gestört wird, während die Rückbestimmung im entsprechenden Fall sehr sicher sein kann.



Fig. 7

Auftreten einer besonders unsicheren Kette.

Jedoch ist es nicht der höhere Wahrscheinlichkeitsgrad, was die Besonderheit des Schlusses in die Vergangenheit ausmacht. Es gibt ja auch Fälle, für die der Vergangenheits-schluß unsicher wird. Sondern die Art der Schlußweise unterscheidet den Rückschluß vom Vorwärtsschluß. Wir wollen, um dies ganz deutlich zu machen, noch eine weitere Struktur betrachten. Die Doppelgabel der Fig. 8 ist der Verkettung nach für Vergangenheit und Zukunft symmetrisch. Aber in den Wahrscheinlichkeitsrelationen ergibt sich Unsymmetrie. Es gilt nämlich für die Doppelgabel:

$$\begin{aligned} A.B &\rightarrow C.D \\ C \vee D &\rightarrow A.B \end{aligned} \quad (6)$$

$A.B$ ist die Gesamtursache, $C.D$ die Gesamtwirkung. Daß man vom Ganzen auf den Teil schließen kann, gilt allerdings für beide Richtungen; dies rührt her von einer Grundeigenschaft der Implikation, die wir schreiben können: $a.b \supset a$. Es ist in den Gesetzen der Wahrscheinlichkeitsimplikation durch 2^* ausgedrückt, d. h. durch die Auflösbarkeit des hinteren „und“, welche für beide Gleichungen (6) gilt. Aber beim Schluß in die Vergangenheit kann man vom Teil zum Ganzen schließen — dies besagt das vordere „oder“ in der zweiten Gleichung — während dies beim Zukunftsschluß unmöglich ist — hier steht vorn ein „und“.

Von dieser Erkenntnis aus gelingt eine Begriffsbestimmung, die wir eingangs schon berührt haben. Wir nannten dort die Zukunft objektiv unbestimmt, im Gegensatz zur Vergangenheit,

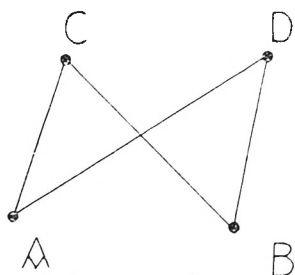


Fig. 8

Doppelgabel.

die objektiv bestimmt ist. Was bedeutet nun objektive Bestimmtheit? Man ist leicht zu folgender Definition geneigt: Ein Zustand ist objektiv bestimmt, wenn die Wahrscheinlichkeit, mit der er subjektiv bestimmt werden kann, beliebig nahe an 1 gesteigert werden kann. Diese Definition hat zunächst den Nachteil, daß sie den Grad der Bestimmbarkeit benutzt, um das Objektive zu definieren. Sie wird aber ganz unhaltbar, wenn man die Annahme fallen läßt, daß eine Grenzfunktion existiert, die einen Weltquerschnitt mit völliger Strenge beschreibt. Denn dann entspricht der Wahrscheinlichkeit 1 gar kein definierter Weltzustand; der Grenzfall ist ausgeartet und kann nicht zur Definition des Objektiven verwandt werden.

Wir können aber auf andere Weise den Begriff „objektiv bestimmt“ gewinnen. Wir werden die Vergangenheit objektiv bestimmt nennen, weil sie aus einer Teilwirkung schon erschlossen werden kann. Denn ein Schluß vom Teil zum Ganzen setzt voraus, daß das Ganze bereits unabhängig feststeht. Wir verfolgen ja mit dem Begriff „objektiv bestimmt“ den Gedanken, daß wir den Zustand nicht mehr ändern können; eben dies bringt die Eigenart des Vergangenheitsschlusses zum Aus-

druck, der ein Bezeugen, nicht ein Bewirken charakterisiert. Ein Bezeugen können schon Teilwirkungen leisten; niemals aber können Teilursachen das Geschehen hervorbringen. Darum kann eine Aussage über die Zukunft erst gewonnen werden, wenn es feststeht, daß alle Teilursachen da sind; aber für den Vergangenheitsschluß sind nicht alle Teilwirkungen notwendig. Es ist charakteristisch, daß wir vergangene Ereignisse registrieren können. Welche Temperatur herrschte vorgestern in diesem Zimmer? Wollten wir dies aus den heute noch irgend wie vorhandenen Wirkungen erschließen, so kommen wir in große Schwierigkeiten. Steht aber ein Registrierthermometer im Zimmer, so ist es leicht, die Antwort zu finden; die Wirkungskette, die an diesem Apparat angreift, können wir zu einem Rückschluß großer Wahrscheinlichkeit verwenden, und die weiteren Wirkungen brauchen wir für den Rückschluß nicht mehr. Eine analoge Einrichtung für die Zukunft ist aber nicht möglich. Wir können die Zukunft nicht registrieren, d. h. eine einzelne Teilkette genügt nicht, um sie zu bestimmen.¹⁾

Die Zukunft müssen wir deshalb „objektiv unbestimmt“ nennen. Denn wenn die Grenzfunktion nicht existiert, so ist die Gesamtheit aller Teilursachen keine definierte Größe. Man kann dann nicht sagen, daß es nur ein Mangel an technischen Mitteln ist, der die Bestimmtheit des zukünftigen Weltzustandes unter die Gewißheitsgrenze drückt; sondern die Unbestimmtheit ist eine objektive Eigenschaft der Kausalstruktur.

So reduziert sich der Unterschied von „objektiv bestimmt“ und „objektiv unbestimmt“ auf einen topologischen Unterschied der Wahrscheinlichkeitsimplikation; der Unterschied der beiden Begriffe „oder“ und „und“ wird nicht nur entscheidend für den Unterschied von Vergangenheit und Zukunft, sondern auch für die Charakterisierung des objektiv Bestimmten im Gegensatz zum Unbestimmten. So sehr auch eine solche Begriffsbestimmung im Anfang befremden mag — wenn man sich einmal von dem Gedanken frei gemacht hat, das objektiv Feststehende durch den Grad der Bestimmbarkeit zu charakterisieren, erscheint sie sehr

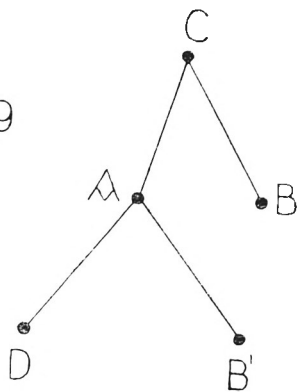
¹⁾ Darum gibt es eine Geschichtswissenschaft nur von der Vergangenheit. Die Chronik, d. i. das Registrieren der Ereignisse, ist das typische Kennzeichen der Geschichte.

viel glücklicher als diese. Denn sie benutzt einen qualitativen Unterschied, und nicht einen quantitativen, zur Kennzeichnung des Objektiven. Sie soll doch schließlich die Tatsache zum Ausdruck bringen, daß die Vergangenheit jedem Wirkungseinfluß entzogen ist; aber das ist ein qualitativer Unterschied, der nicht durch einen Wahrscheinlichkeitsgrad getroffen werden kann.

Man könnte den Gedanken, daß der Wirkungseinfluß sich nur zeitlich vorwärts ausbreiten kann, auch folgendermaßen zu formulieren suchen. Ist A die Ursache von C , und bringt man in A eine kleine Änderung an, so wird auch in C eine kleine Änderung auftreten. Wenn man aber in C eine kleine Änderung anbringt, so entsteht in A keine Änderung. Aber die Formulierung mit Hilfe der willkürlich angebrachten Änderung ist anfechtbar, weil es vom Standpunkt des Determinismus gar nicht möglich ist, eine Änderung willkürlich anzubringen. Dieser Fehler wird vermieden, wenn man die Wahrscheinlichkeitsimplikation benutzt. Eine Änderung in C

bedeutet eben, daß dort noch eine zweite, nicht von A kommende Kausalkette eintrifft, so daß die Spitzgabel (Fig. 9) entsteht; und daß diese zusätzliche von B kommende Kette keinen Einfluß auf A hat, drückt sich in dem Fehlen einer Wahrscheinlichkeitsimplikation zwischen B und A aus. Greift dagegen die zusätzliche Kette, von B' kommend, in A an (Fig. 9), so gilt wegen $C \rightarrow A$ und $A \rightarrow B'$ nach

Fig. 9



Spitzgabeln als Erklärung für die Einseitigkeit des Wirkungseinflusses.

9* $C \rightarrow B'$, d. h. es ist die Wirkung von B' in C zu beobachten. Der Begriff der Wahrscheinlichkeitsimplikation erlaubt also die einwandfreie Formulierung der Tatsache, daß die Wirkung sich nur zeitlich vorwärts, nicht rückwärts ausbreitet.

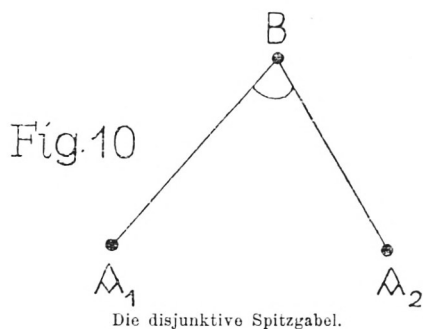
III. Zusammenhang von Vor- und Rückwahrscheinlichkeit.

Wir haben im Vorangehenden angenommen, daß sowohl eine Wahrscheinlichkeit für die Richtung „zeitlich vorwärts“ als auch

für die Richtung „zeitlich rückwärts“ besteht. Wir wollen zeigen, welche Voraussetzung in diesen Annahmen enthalten ist, und wie sich die Rückwahrscheinlichkeit aus der Vorwahrscheinlichkeit berechnen läßt. Dazu müssen wir in das Strukturbild noch eine andere Art der Verknüpfung eintragen, die Verknüpfung mit möglichen Kausalketten.

1. Die disjunktive Spitzgabel. Es sei B eine Wirkung, die sowohl bei Auftreten der Ursache A_1 als auch bei A_2 entsteht; aber B soll nicht durch ein Zusammenwirken der beiden Ursachen entstehen, sondern gerade nur dann eintreten, wenn nur eine der Ketten $A_1 B$ oder $A_2 B$ vorliegt. Dies unterscheidet den Fall von der konjunktiven Spitzgabel des vorigen Abschnitts. Zur Kennzeichnung der disjunktiven Eigenschaft setzen wir einen Winkelbogen in die Figur.

Für die disjunktive Eigenschaft wollen wir eine besondere Schreibweise einführen. Wir dürfen nicht sagen, daß die Kombination $A_1 A_2$ mit dem Ein-



treten von B unvereinbar ist, denn A_2 (bzw. A_1) kann, da es B nur mit Wahrscheinlichkeit impliziert, eine andere Wirkung Q_i haben, während A_1 gerade B liefert. Nur wenn keine weitere Wirkung Q_i von A_1 oder A_2 vorliegt, ist B mit der Kombination $A_1 A_2$ unvereinbar; denn da wir an-

nehmen, daß eine Kausalkette niemals endet, müssen dann beide Ketten $A_1 B$ und $A_2 B$ vorliegen, und dies ist nach Voraussetzung mit B unvereinbar. Wir schreiben deshalb:

$$B \Rightarrow A_1 \wedge A_2 = \begin{cases} B.Pl(\bar{Q}_i) \Rightarrow A_1 \cdot A_2 & [0] \\ B.Pl(Q_i) \Rightarrow A_1 & [p] \\ B.Pl(Q_i) \Rightarrow A_2 & [q] \\ Q_i \Rightarrow A_1 \vee A_2 \end{cases} \quad Df. \quad (7)$$

Hier bedeutet $Pl(\bar{Q}_i)$ das logische Produkt aller möglichen Q_i , also

$$Pl(\bar{Q}_i) = \bar{Q}_1 \cdot \bar{Q}_2 \dots \bar{Q}_n \quad Df. \quad (8)$$

Der Buchstabe bzw. die Zahl in der [] bedeutet das Maß der betreffenden Wahrscheinlichkeitsimplikation; wesentlich in der gegebenen Definition ist, daß dieses Maß in der ersten Zeile = 0 wird. Wir lesen den in (7) links stehenden Ausdruck, der durch die rechte Seite definiert wird, als „ B allein bestimmt entweder A_1 oder A_2 “. Das Zeichen $\wedge$ bedeutet also das „ausschließende oder“; aber man muß beachten, daß wir nicht dieses Zeichen, sondern nur den Ausdruck $B \rightarrow A_1 \wedge A_2$ definiert haben, und daß dieser Ausdruck noch die Bedeutung von „ B allein“ enthält.

Wir können nun die Relationen der disjunktiven Spitzgabel schreiben:

$$A_1 \vee A_2 \rightarrow B \quad B \rightarrow A_1 \cdot A_2 \quad B \rightarrow A_1 \wedge A_2. \quad (9)$$

Hieraus folgt mit (7), 2*, 3*, 5*, 9*

$$\begin{array}{ccccccc} A_1 \rightarrow B & A_2 \rightarrow B & B \rightarrow A_1 & B \rightarrow A_2 & & & \\ & & A_1 \rightarrow A_2 & A_2 \rightarrow A_1 & & & \end{array} \quad (10)$$

Hier ist zunächst das vordere „oder“ in dem ersten der Ausdrücke (9) auffallend, da es sich in diesem Ausdruck um einen Zukunftsschluß handelt. Dieses „oder“ kann nur deshalb eintreten, weil der Ausdruck $B \rightarrow A_1 \wedge A_2$ hinzutritt, also der disjunktive Fall vorliegt. Sodann ist die Folgerung $A_1 \rightarrow A_2$ und $A_2 \rightarrow A_1$ auffallend, welche ein Zahlenverhältnis zwischen Ereignissen besagt, die nicht durch eine gemeinsame Ursache, sondern durch eine Wirkung miteinander verbunden sind. Freilich ist dies nicht eine gemeinsame Wirkung, sondern eine gleiche Wirkung, und zwar eine mögliche gleiche Wirkung. Wir wollen, um die Richtigkeit unserer Schlüsse (und damit auch unserer Gesetze der Wahrscheinlichkeitsimplikation) zu zeigen, diesen Fall genauer verfolgen.

Sind alle Aussagen (9) und (10) richtig? Insbesondere die erste Aussage (9) und die letzten beiden Aussagen (10) erscheinen zweifelhaft. Aber irgendwelche Wahrscheinlichkeitsimplikationen müssen hier bestehen, und zwar folgt aus dem Sinn des Problems:

$$A_1 \cdot \bar{A}_2 \rightarrow B \quad [u] \quad (11)$$

$$\bar{A}_1 \cdot A_2 \rightarrow B \quad [v] \quad (12)$$

$$B \rightarrow A_1 \cdot A_2 \quad [u'] \quad (13)$$

$$B \rightarrow \bar{A}_1 \cdot A_2 \quad [v'] \quad (14)$$

Die ersten beiden dieser Gleichungen besagen die Definition des Problems, die letzten beiden die Annahme der entsprechenden beiden Rückwahrscheinlichkeiten. Es muß aber noch eine weitere Gleichung gelten:

$$A_1 \cdot A_2 \Rightarrow B \quad [s] \quad (15)$$

Denn falls A_1 und A_2 da sind, können wir B als Spitze einer konjunktiven Spitzgabel deuten, deren Ketten die Wahrscheinlichkeit u und $1-v$ bzw. $1-u$ und v haben; es gilt also

$$s = u(1-v) + v(1-u) = u + v - 2uv \quad (16)$$

Wir wollen zeigen, daß aus den 5 Gleichungen (11)–(15) die Beziehungen (9) und (10) sämtlich folgen; und zwar wollen wir dies nicht mit Hilfe unserer Gesetze der Wahrscheinlichkeitsimplikation zeigen — damit wäre es sofort bewiesen — sondern durch Ausrechnen. Wir wollen alle möglichen Kombinationen der 3 Ereignisse A_1 , A_2 , B in ihrem Eintreffen und Nichteintreffen verfolgen. Dabei wollen wir der Einfachheit halber annehmen, daß außer A_1 und A_2 keine weiteren Ursachen für B möglich sind; dies bedeutet keine Einschränkung unserer Behauptungen, sondern verringert nur die Zahl der Unbekannten und der Gleichungen. In der folgenden Tabelle sind die möglichen Kombinationen mit ihren Häufigkeiten, die man sich experimentell beobachtet denken möge, hingeschrieben:

$$\begin{array}{ll} A_1 \cdot A_2 \cdot B & n_1 \\ A_1 \cdot A_2 \cdot \bar{B} & n_2 \\ \bar{A}_1 \cdot A_2 \cdot B & n_3 \\ \bar{A}_1 \cdot A_2 \cdot \bar{B} & n_4 \\ A_1 \cdot \bar{A}_2 \cdot B & n_5 \\ A_1 \cdot \bar{A}_2 \cdot \bar{B} & n_6 \end{array} \quad (17)$$

Wegen der genannten vereinfachenden Voraussetzung existiert eine Kombination $\bar{A}_1 \cdot A_2 \cdot B$ nicht; die Kombination $\bar{A}_1 \cdot A_2 \cdot \bar{B}$ brauchen wir nicht zu berücksichtigen, da sie in keiner der Beziehungen (9) und (10) mitgezählt ist. Es entsteht nun die Frage: werden die einmal beobachteten Häufigkeiten $n_1 \dots n_6$ in ihrem Verhältnis ungeändert erhalten, wenn man die Zählung über eine größere Zahl von Ereignissen erstreckt?

Durch die 5 Beziehungen (11)–(15) sind die 5 Gleichungen gegeben:

$$\frac{n_1}{n_1 + n_2} = u \quad (18)$$

$$\frac{n_3}{n_3 + n_4} = v \quad (19)$$

$$\frac{n_1}{n_1 + n_3 + n_5} = u' \quad (20)$$

$$\frac{n_3}{n_1 + n_3 + n_5} = v' \quad (21)$$

$$\frac{n_5}{n_5 + n_6} = s \quad (22)$$

Diese 5 Gleichungen sind voneinander unabhängig. (18), (19), (22) bedeuten nur den Zusammenhang zwischen n_1 und n_2 , n_3 und n_4 , n_5 und n_6 , und sind unabhängig sowohl voneinander als von (20) und (21). Diese letzteren beiden sind aber ebenfalls voneinander unabhängig.

Es sind also mit (18)–(22) die Verhältnisse der 6 Unbekannten $n_1 \dots n_6$ festgelegt, und darum müssen diese Verhältnisse konstant bleiben, wenn die Beziehungen (11)–(15) gelten. Dann muß aber auch jede andere Wahrscheinlichkeitsbeziehung zwischen den Größen A_1 , A_2 , B gelten, denn sie wird durch Auszählen der Häufigkeiten $n_1 \dots n_6$ gewonnen. Z. B. wird die Wahrscheinlichkeit für $A_1 \rightarrow A_2$

$$t = \frac{n_5 + n_6}{n_1 + n_2 + n_5 + n_6} \quad (23)$$

und dies muß konstant bleiben, wenn die Verhältnisse der $n_1 \dots n_6$ konstant bleiben. Damit ist bewiesen, daß die Beziehungen (9) und (10) die Folge von (11)–(15) sind, ohne daß für den Beweis die Gesetze der Wahrscheinlichkeitsimplikation benutzt werden.¹⁾

Es sei noch der Fall besprochen, daß sowohl A_1 als auch A_2 keine andere Wirkung haben können als B , also $A_1 \rightarrow B$ [1],

¹⁾ Der Beweis ließe sich ebenso führen, wenn man an Stelle von (13) und (14) die Relationen $B \rightarrow A_1$ und $B \rightarrow A_2$ oder auch an Stelle von (11) und (12) die Relationen $A_1 \rightarrow B$ und $A_2 \rightarrow B$ benutzen würde.

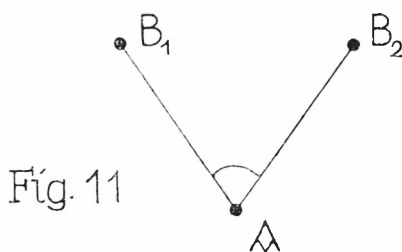
$A_2 \rightarrow B$ [1]. Dann ist $A_1 \cdot A_2 \cdot B$ unmöglich und $n_5 = 0$. Hier werden aber (20) und (21) voneinander abhängig, weil dann $u' + v' = 1$ wird. Es sind also 4 Gleichungen für 5 Unbekannte vorhanden und wieder nur die Verhältnisse der $n_1 n_2 n_3 n_4 n_5$ festgelegt.

Wir können jetzt die Konsequenzen der Beziehungen (9) und (10) verfolgen. Mit $A_1 \rightarrow A_2$ und $A_2 \rightarrow A_1$ ist gesagt, daß die beiden Ereignisse A_1 und A_2 in ihrem Auftreten eine regelmäßige Häufigkeit befolgen. Wenn man alle Ereignisse der Welt durchsieht, so findet man, daß die Häufigkeit von A_1 und A_2 ein konstantes Verhältnis zeigt. Nicht nur die gemeinsame Ursache, sondern auch die mögliche gleiche Wirkung stellt eine Wahrscheinlichkeits- und Häufigkeitsbeziehung zwischen Ereignissen her.

Dieses Resultat folgt aus der Existenz einer Vor- und Rückwahrscheinlichkeit. Würde nur die Rückwahrscheinlichkeit gelten, so würde eine regelmäßige Häufigkeit nur zwischen BA_1 und BA_2 gelten, d. h. man dürfte nur die Fälle zählen, in denen A_1 oder A_2 von B begleitet sind.

2. Die disjunktive Satteltgabel. Hier gelten dieselben Relationen wie bei der disjunktiven Spitzgabel; die disjunktive Gabel liefert also keine Auszeichnung einer Richtung

$$B_1 \vee B_2 \rightarrow A \quad A \rightarrow B_1 \cdot B_2 \quad A \rightarrow B_1 \wedge B_2 \quad (24)$$



Die disjunktive Satteltgabel.

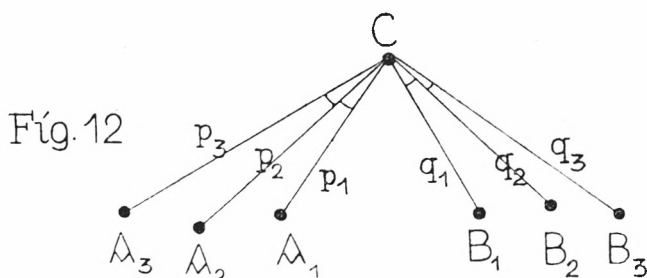
Die Aussage $A \rightarrow B_1 \wedge B_2$ ist dabei gerade so definiert wie in (7) angegeben, nur daß hier unter Q_i die weiteren möglichen Ursachen von B_1 oder B_2 zu verstehen sind. — Die Relationen sind hier dieselben wie bei der konjunktiven Satteltgabel, das Hinzutreten der disjunktiven Beziehung bedeutet also keine Änderung.

Das Resultat über die Häufigkeit von Ereignissen, das dem obigen entspricht, heißt hier: die gemeinsame Ursache stellt eine Häufigkeitsbeziehung zwischen Ereignissen auch dann her, wenn immer nur eines der Ereignisse Wirkung der betreffenden Ursache sein kann.

Beide Aussagen zusammen formulieren wir als das

Verteilungsgesetz für Ereignisse in der Welt: Solche Ereignisse, die auf eine gleiche oder gemeinsame Ursache zurückgeführt werden können, oder eine gleiche Wirkung haben können, zeigen in ihrem Auftreten in der Welt eine regelmäßige gegenseitige Häufigkeit.

3. Die konjunktive Spitzgabel. Wir gehen jetzt dazu über, für einige Fälle die Rückwahrscheinlichkeit aus der Vorwahrscheinlichkeit quantitativ zu berechnen. Wir beginnen mit der konjunktiven Spitzgabel. In der Fig. 12 sind die Vorwahrscheinlichkeiten eingetragen; sie heißen p_i und q_i , während die entsprechenden Rückwahrscheinlichkeiten p'_i und q'_i heißen sollen.



Spitzgabel mit eingezeichneten andern möglichen Ursachen.

Für das Maß der Wahrscheinlichkeit in $A \rightarrow B$, welches wir früher in [] der Aussage hinzugefügt haben, wollen wir die Bezeichnung $W(A \rightarrow B)$ einführen. Dann wird, da A_1 und B_1 unabhängige Ereignisse sein sollen:

$$W(A_1.B_1 \rightarrow C) = r_1 = p_1 \cdot q_1 \quad (25)$$

Hier bedeuten die Größen p_i und q_i nicht etwa die Wahrscheinlichkeiten $W(A_i \rightarrow C)$ und $W(B_i \rightarrow C)$, da diese Implikationen nach (3) nicht existieren. Sondern sie bedeuten nur die Wahrscheinlichkeit, daß die von A_i bzw. B_i herkommende Teilwirkung in C eintritt. Daher ergibt erst ihr Produkt die Wahrscheinlichkeit r_i .

Um aus dieser Vorwahrscheinlichkeit auf die Rückwahrscheinlichkeit schließen zu können, müssen wir A_i und B_i ($i > 1$), die andern möglichen Ursachen für C , hinzuziehen. Und zwar müssen

wir über die sämtlichen möglichen Ursachen für C eine Annahme machen, etwa eine der beiden folgenden Annahmen:

Annahme H . Irgend zwei Ereignisse $A_i B_k$ können zusammen C liefern.

Dann ist

$$W(A_i B_k \rightarrow C) = p_i \cdot q_k \quad (26)$$

Annahme J . Zu jedem A_i muß ein bestimmtes B_i hinzutreten, damit C entsteht.

Dann ist

$$\begin{aligned} W(A_i B_i \rightarrow C) &= p_i \cdot q_i \\ W(A_i B_k \rightarrow C) &= 0 \quad i \neq k \end{aligned} \quad (27)$$

Wir fragen jetzt: wie groß ist

$$W(C \rightarrow A_1) = p'_1 \quad W(C \rightarrow B_1) = q'_1 \quad W(C \rightarrow A_1 B_1) = r'_1$$

Für die Berechnung benutzen wir die Bayessche Regel.¹⁾ Diese lautet: Sind $X_1 \dots X_n$ alle möglichen Ursachen für Y , und ist

$$W(X_i \rightarrow Y) = z_i \quad W(X_i) = a_i$$

so ist

$$W(Y \rightarrow X_i) = z'_i = \frac{a_i z_i}{\sum_1^n a_k z_k} \quad (28)$$

$W(X_i)$ ist die sogenannte „apriorische Wahrscheinlichkeit“ für X_i , d. h. die relative Wahrscheinlichkeit der X_i gegeneinander in ihrem Auftreten in der Welt. Diese Größen a_i können auch sämtlich gleich groß werden, dann wird²⁾

$$z'_i = \frac{z_i}{\sum_1^n z_k} \quad (29)$$

¹⁾ Vgl. jedes Lehrbuch der Wahrscheinlichkeitsrechnung, etwa Czuber, 1908, S. 175.

²⁾ Wie man aus (28) und (29) sieht, wird $\sum_1^n z'_i = 1$; die z'_i heißen deshalb „verbundene Wahrscheinlichkeiten.“ Dagegen sind die z_i „unverbundene Wahrscheinlichkeiten“, d. h. $\sum_1^n z_i \geq 1$.

Ohne die α_i ist jedoch das Problem nicht vollständig definiert und p'_i nicht berechenbar. Die α_i sind offenbar die relativen Wahrscheinlichkeiten unseres Verteilungsgesetzes; wir wollen sie deshalb Verteilungswahrscheinlichkeiten nennen, da die Verwendung des Begriffs „apriori“ in diesem Zusammenhang irreführen kann. Nur weil sie existieren, ist die Rückwahrscheinlichkeit aus der Vorwahrscheinlichkeit berechenbar. In den Schulbeispielen der Bayesschen Regel werden die α_i gewöhnlich dadurch begründet, daß man die X_i auf eine gemeinsame Ursache X zurückführt, welche die X_i disjunktiv erzeugt. Diese Begründung ist jedoch nicht notwendig. Wenn eine Rückwahrscheinlichkeit existiert, so müssen die α_i existieren; dies genügt uns als Begründung.

Mit der Bayesschen Regel nach (28) kommen wir jetzt auf unser Problem zurück. Sei $W(A_i) = \alpha_i$ und $W(B_i) = \beta_i$. Die Anzahl der möglichen Ursachen A_i sei m , die der B_i sei n . Wir können jede Kombination $A_i B_k$ als Einzelereignis auffassen, welches die apriorische Wahrscheinlichkeit $\alpha_i \beta_k$ hat und mit der Wahrscheinlichkeit $p_i q_k$ die Wirkung C erzeugt. Dann wird nach (28), wenn wir Annahme H zugrunde legen:

$$\begin{aligned} p_i &= \frac{\alpha_i p_1 \cdot \sum_1^n \beta_i q_i}{\sum_1^m \sum_1^n \alpha_i p_i \beta_k q_k} = \frac{\alpha_i p_1}{\sum_1^m \alpha_i p_i} \\ q'_1 &= \frac{\beta_1 q_1 \sum_1^m \alpha_i p_i}{\sum_1^m \sum_1^n \alpha_i p_i \beta_k q_k} = \frac{\beta_1 q_1}{\sum_1^n \beta_i q_i} \\ r'_1 &= \frac{\alpha_1 p_1 \beta_1 q_1}{\sum_1^m \sum_1^n \alpha_i p_i \beta_k q_k} = p'_1 \cdot q'_1 \end{aligned} \quad (30)$$

Mit Annahme J wird dagegen (hier muß $m = n$ sein)

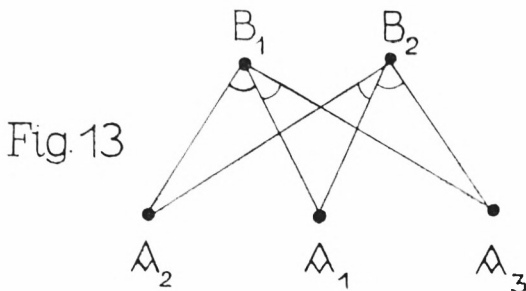
$$p'_i = q'_1 = r'_1 = \frac{\alpha_1 p_1 \beta_1 q_1}{\sum_1^m \alpha_i p_i \beta_i q_i} \quad (31)$$

Bei Annahme H berechnet sich also die Rückwahrscheinlichkeit r'_1 aus den einzelnen Rückwahrscheinlichkeiten p'_i und q'_i ge-

rade so wie sich die entsprechende Vorwahrscheinlichkeit r_i aus den einzelnen Vorwahrscheinlichkeiten p_i und q_i berechnet. Bei Annahme J dagegen wird die Berechnung für die Rückwahrscheinlichkeit anders, und zwar werden hier alle drei Rückwahrscheinlichkeiten gleich. Ob Annahme H oder J der Wirklichkeit entspricht, hängt natürlich von den besonderen Bedingungen des Problems ab; im allgemeinen wird wohl ein aus beiden Annahmen gemischter Fall vorliegen. Dann berechnet sich für r_i ein entsprechender Zwischenwert.

4. Die konjunktive Satteltabel. Auch hier müssen wir, um die Rückwahrscheinlichkeit aus der Vorwahrscheinlichkeit berechnen zu können, noch eine Annahme hinzunehmen. Wir wählen zunächst die

Annahme K : B_1 und B_2 haben notwendig eine gemeinsame Ursache.



Satteltabel mit eingezeichneten andern möglichen Ursachen nach Annahme K .

Seien $A_1 \dots A_m$ (Fig. 13) die möglichen gemeinsamen Ursachen; ihre Verteilungswahrscheinlichkeit sei $W(A_i) = a_i$. Wir führen noch die Bezeichnungen ein:

$$\begin{aligned} W(A_i \rightarrow B_1) &= p_i & W(B_1 \rightarrow A_i) &= p'_i \\ W(A_i \rightarrow B_2) &= q_i & W(B_2 \rightarrow A_i) &= q'_i \\ W(A_i \rightarrow B_1 B_2) &= r_i & W(B_1 B_2 \rightarrow A_i) &= r'_i \end{aligned}$$

Dann wird

$$p'_i = \frac{a_i p_i}{\sum_1^m a_k p_k} \quad q'_i = \frac{a_i q_i}{\sum_1^m a_k q_k} \quad (32)$$

$$r'_i = \frac{\alpha_i p_i q_i}{\sum_1^m \alpha_k p_k q_k} = \frac{p'_i \cdot q'_i}{\alpha_i \sum_1^m \frac{1}{\alpha_k} p'_k q'_k} \quad (32)$$

Für die Diskussion dieser Formeln nehmen wir den einfachen Fall an, daß die α_i alle gleich groß werden. Dann reduzieren sich die Formeln (32) auf

$$p'_i = \frac{p_i}{\sum_1^m p_k} \quad q'_i = \frac{q_i}{\sum_1^m q_k} \quad r'_i = \frac{p_i q_i}{\sum_1^m p'_k q'_k} \quad (33)$$

Die Größen p'_i , q'_i , r'_i sind als Wahrscheinlichkeiten alle < 1 , wie man auch leicht sieht. Es wird aber außerdem

$$r'_i > p'_i \cdot q'_i \quad (34)$$

denn

$$\sum_1^m p'_k q'_k < \sum_1^m \sum_1^m p'_k q'_i = \sum_1^m p'_k \cdot \sum_1^m q'_i = 1$$

Würde man die Satteltabel nach oben klappen, so daß eine Spitzgabel entstände, so würde sich, wenn p'_i , q'_i , r'_i jetzt die entsprechenden Vorwahrscheinlichkeiten bedeuten, $r'_i = p'_i \cdot q'_i$ ergeben.¹⁾ Man erkennt: der Vergangenheitsschluß von zwei Ereignissen auf eines liefert eine größere Wahrscheinlichkeit als der entsprechende Zukunftsschluß bei gleichen Wahrscheinlichkeiten. So drückt sich auch in metrischer Beziehung die Besonderheit des Vergangenheitsschlusses aus, als einer Aussage über ein Bezeugen, im Gegensatz zum Bewirken; schon jede Einzelwirkung erlaubt den Rückschluß auf die Ursache mit gewisser Wahrscheinlichkeit, und eine hinzukommende Wirkung bedeutet nur eine Bestätigung, nicht eine notwendige Bedingung für den Rückschluß.

¹⁾ Wenn die α_i nicht gleich groß sind, muß (34) nicht erfüllt sein. Aber dann würde auch $p'_i q'_i$ nicht die Bedeutung der Vorwahrscheinlichkeit bei der Spitzgabel haben, weil dann noch die apriorische Wahrscheinlichkeit von A_i hinzutritt; p'_i und q'_i sind dann keine unabhängigen Wahrscheinlichkeiten. — Man beachte ferner: wir vergleichen hier nicht die Rückwahrscheinlichkeit mit der Vorwahrscheinlichkeit desselben Falles, also nicht r'_i mit r_i , sondern mit der Vorwahrscheinlichkeit eines entsprechenden Falles, in dem die p'_i und q'_i Vorwahrscheinlichkeiten (und zwar Wahrscheinlichkeiten der Teilwirkung im Sinne unserer auf (25) folgenden Bemerkung) bedeuten; also wir vergleichen r'_i mit $p'_i q'_i$.

Freilich kann diese Bestätigung auch negativen Charakter haben, wenn die zweite Wirkung die vermutete Ursache gerade „unwahrscheinlich macht“. Um dies zu untersuchen, vergleichen wir r'_i mit der Einzelwahrscheinlichkeit p'_i . Es wird

$$r'_i = p'_i \cdot f \quad f = \frac{q'_i}{\sum_1^m p'_k q'_k} = \frac{q'_i \sum_1^m p_k}{\sum_1^m p_k q_k} \quad (35)$$

Je nachdem $f \leq 1$, wird $r'_i \leq p'_i$. Also wird

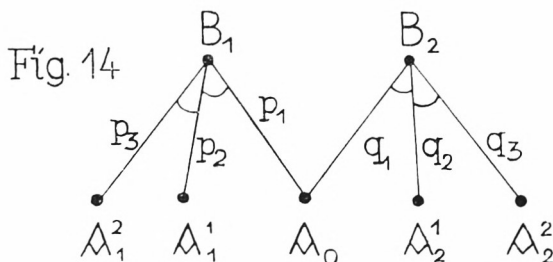
$$r'_i \leq p'_i \text{ wenn } q_i \sum_1^m p_k \leq \sum_1^m p_k q_k \quad (36)$$

$$r'_i \leq q'_i \text{ wenn } p_i \sum_1^m q_k \leq \sum_1^m p_k q_k$$

Sind etwa alle q_k gleich groß, so wird $r'_i = p'_i$, d. h. das Hinzutreten des Ereignisses B_2 bedeutet dann weder eine Vermehrung noch eine Verminderung der Wahrscheinlichkeit, die sich aus B_1 allein für A_i berechnen läßt. Dann sind eben für B_2 alle Ursachen A_i gleichwahrscheinlich. Ist q_i das größte unter allen q_k , so wird $r'_i > p'_i$; ist q_i das kleinste, so wird $r'_i < p'_i$. Je nachdem also B_2 die Ursache A_i „wahrscheinlich macht“ oder „unwahrscheinlich macht“, tritt eine Vermehrung oder Verminderung der Wahrscheinlichkeit ein. Die Bedingung (36) besagt, daß dieses „wahrscheinlich machen“ dann vorliegt, wenn q_i einen mit Hilfe der p_k gebildeten Mittelwert aus den q_k überschreitet.

In den Formeln (32) und (33) ist noch zu beachten, daß sich die zusammengesetzte Rückwahrscheinlichkeit r'_i aus den Rückwahrscheinlichkeiten p'_i und q'_i der Einzelereignisse und den Verteilungswahrscheinlichkeiten α_i ausdrücken läßt, ohne daß die Vorwahrscheinlichkeiten p_i und q_i eingehen. Dies ist eine Besonderheit, die nicht für alle Fälle gilt; sie beruht hier auf der Annahme K . Wir wollen jetzt noch einen andern Fall betrachten, in dem diese Besonderheit ebenfalls gilt, der aber auf einer andern Annahme beruht.

Wir wollen jetzt nicht mehr annehmen, daß die Ereignisse B_1 und B_2 aus einer einzigen Ursache A_0 erklärt werden müssen, sondern für jedes von beiden getrennte Ursachen A_1^i und A_2^k zulassen (Fig. 14). Wir suchen aber gerade die Wahrscheinlich-



Sattelgabel mit eingezeichneten andern möglichen Ursachen, die getrennt sind,

keit, daß die gemeinsame Ursache A_0 vorliegt. Um das Problem zu vereinfachen, wollen wir aber noch eine Annahme hinzufügen:

Annahme L . Wenn A_0 vorliegt, so sind alle $A_1^1 \dots$ und $A_1^m A_2^1 \dots A_2^n$ ausgeschlossen. $A_0 \supset \bar{A}_1^i \cdot \bar{A}_2^k$.

In Wirklichkeit wird ja eine derartige Annahme, wie alle vorherigen, nicht mit völliger Sicherheit gelten; d. h. die strenge Implikation in Annahme L ist nur eine Wahrscheinlichkeitsimplikation von sehr hohem Grade. Aber für praktische Fälle wird es häufig erlaubt sein, diesen hohen Grad der Wahrscheinlichkeit als Gewißheit zu betrachten. So dürfte die Annahme L für viele Fälle des Indizienbeweises zutreffen, in denen es sich darum handelt, die Indizien, die jedes für sich unabhängige Ursachen zulassen, auf eine gemeinsame Ursache zurückzuführen. Für die Rechnung wollen wir außerdem noch die Annahme machen, daß die a_i alle gleich sind; sonst wird das Resultat unübersichtlich. Es wird (Bezeichnungen wie bei der vorigen Rechnung):

$$\begin{aligned}
 p'_0 &= \frac{p_0}{\sum_0^m p_k} & q'_0 &= \frac{q_0}{\sum_0^n q_k} \\
 r'_0 &= \frac{p_0 q_0}{p_0 q_0 + \sum_1^m \sum_1^n p_i q_k} = \frac{p_0 q_0}{\sum_0^m \sum_0^n p_i q_k - p_0 \sum_0^n q_k - q_0 \sum_0^m p_k + 2 p_0 q_0} \\
 &= \frac{p'_0 q'_0}{1 - p'_0 - q'_0 + 2 p'_0 q'_0} = \frac{p'_0 q'_0}{(1 - p'_0)(1 - q'_0) + p'_0 q'_0} \quad (37)
 \end{aligned}$$

Auch hier ist also r'_0 allein durch die Rückwahrscheinlichkeiten ausdrückbar, und es gehen sogar nur die Rückwahrscheinlichkeiten p'_0 und q'_0 ein, während die andern Rückwahrscheinlich-

keiten fortfallen. Wieder ist natürlich $r'_0 < 1$. Es ist aber auch wieder $r'_0 > p'_0 q'_0$, denn der Nenner in (37) ist < 1 . Dies ergibt sich, wenn man beachtet, daß $0 < p'_0 < 1$ und $0 < q'_0 < 1$. Setzt man $p'_0 = 1 - \delta$, $q'_0 = 1 - \eta$, so wird der Nenner $= \delta \eta + p'_0 q'_0 < (p'_0 + \delta)(q'_0 + \eta) = 1$. Wieder berechnet sich also die Rückwahrscheinlichkeit größer als die entsprechende Vorwahrscheinlichkeit.

Wir fragen weiter nach dem Fall, wenn $r'_0 > p'_0$. Dazu muß

$$\frac{q'_0}{1 - p'_0 - q'_0 + 2p'_0 q'_0} > 1$$

sein. Dies führt auf

$$q'_0 > \frac{1}{2} \quad (38)$$

Entsprechend ist $r'_0 > q'_0$, wenn $p'_0 > \frac{1}{2}$. Wir finden also: das zweite Ereignis B_2 verstärkt die aus B_1 berechnete Wahrscheinlichkeit für A_0 , wenn es allein A_0 mit einer größeren Wahrscheinlichkeit als $\frac{1}{2}$ impliziert.¹⁾

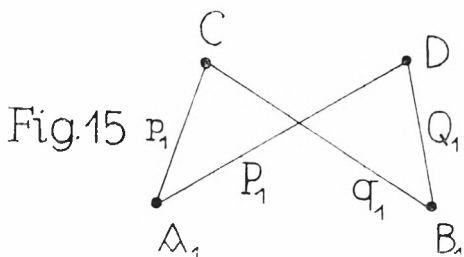


Fig.15 Doppelgabel mit eingezeichneten Vorwahrscheinlichkeiten.

5. Die konjunktive Doppelgabel. Wir betrachten endlich den Fall einer doppelten Verknüpfung. Zwei unabhängige Ursachen A_1 und B_1 haben zwei gemeinsame Wirkungen C und D . Wir wollen diesen Fall auf die schon behandelten Fälle zurückführen.

Wir können den Schluß auf zwei Wegen ausführen, die wir schematisch so schreiben:

a) $C \rightarrow A_1 B_1 [r'_1]$	b) $CD \rightarrow A_1 [s'_1]$
$D \rightarrow A_1 B_1 [r'_1]$	$CD \rightarrow B_1 [s'_1]$
$CD \rightarrow A_1 B_1 [w'_1]$	$CD \rightarrow A_1 B_1 [w'_1]$

Auf dem Weg a) bedeuten die ersten beiden Zeilen einen Vergangenheitsschluß in der konjunktiven Spitzgabel; dann wird $A_1 B_1$ als ein Ereignis aufgefaßt, und in der dritten Zeile ein

¹⁾ Hätten wir die a_i nicht gleich groß gesetzt, so würde hier nicht gerade $\frac{1}{2}$ stehen, sondern eine Bildung aus den a_i .

Vergangenheitsschluß in der Sattलगabel ausgeführt. Auf dem Weg b) bedeuten die ersten beiden Zeilen einen Vergangenheitschluß in der Sattलगabel, dann wird CD als ein Ereignis aufgefaßt, und in der dritten Zeile ein Vergangenheitschluß in der Spitzgabel ausgeführt. Beide Wege müssen zum gleichen Ziele führen. Aber es werden dabei Annahmen gemacht werden müssen, wie wir sie bei den betreffenden Einzelgabeln gemacht haben, damit der Schluß erst definiert ist. Wir wollen für die Sattलगabel die Annahme K voraussetzen, für die Spitzgabel nacheinander Annahme J und H . Der übersichtlicheren Rechnung wegen wollen wir auch hier wieder die Verteilungswahrscheinlichkeiten α_i alle gleich groß setzen.

Wir wählen den Weg a). Es wird zunächst mit Ausnahme J nach (31):

$$w(C \rightarrow A_i B_i) = r'_i = \frac{p_i q_i}{\sum_1^m p_i q_i} \quad w(D \rightarrow A_i B_i) = R'_i = \frac{P_i Q_i}{\sum_1^m P_i Q_i}$$

Mit Annahme K wird nun nach (33)

$$w_i = \frac{r'_i R'_i}{\sum_1^m r'_i R'_i} = \frac{p_i q_i P_i Q_i}{\sum_1^m p_i q_i P_i Q_i} \quad (39)$$

Dagegen wird die Vorwahrscheinlichkeit w_1 für $A_1 B_1 \rightarrow CD$

$$w_1 = p_1 q_1 P_1 Q_1 \quad (40)$$

Definieren wir Größen p'_i, q'_i, P'_i, Q'_i durch

$$p'_i = \frac{p_i}{\sum_1^m p_k} \quad q'_i = \frac{q_i}{\sum_1^n q_k} \quad P'_i = \frac{P_i}{\sum_1^m P_k} \quad Q'_i = \frac{Q_i}{\sum_1^n Q_k} \quad (41)$$

die wir in diesem Falle aber nicht als Rückwahrscheinlichkeiten entsprechend den ungestrichenen Größen deuten dürfen, weil diese Rückwahrscheinlichkeiten wegen Annahme J nach (31) gleich r'_i bzw. R'_i werden, so läßt sich (39) schreiben:

$$w_i = \frac{p'_i q'_i P'_i Q'_i}{\sum_1^m p'_i q'_i P'_i Q'_i} \quad (42)$$

Jetzt wollen wir denselben Weg gehen, aber nicht Annahme J , sondern Annahme H benutzen. Dann wird nach (30)

$$\begin{aligned} w(C \Rightarrow A_i B_k) &= r'_{ik} = \frac{p'_i q'_k}{\sum_1^m \sum_1^n p_i q_k} = p'_i q'_k \\ w(D \Rightarrow A_i B_k) &= R'_{ik} = \frac{P_i Q_k}{\sum_1^m \sum_1^n P_i Q_k} = P_i Q_k \end{aligned} \quad (43)$$

wo wir nun den p'_i , q'_i , P_i , Q_i die Deutung als Rückwahrscheinlichkeiten entsprechend den ungestrichenen Größen beilegen dürfen, da Annahme H dies zuläßt. Weiter wird mit Annahme K nach (33):

$$w_1 = \frac{p'_1 q'_1 P_1 Q_1}{\sum_1^m \sum_1^n p'_i q'_k P_i Q_k} = \frac{p_1 q_1 P_1 Q_1}{\sum_1^m \sum_1^n p_i q_k P_i Q_k} \quad (44)$$

Es ergibt sich also ein Resultat, das von (42) bzw. (39) durch die Doppelsumme im Nenner verschieden ist.

Wenn man den Weg b) benutzt, so läßt sich Annahme J nicht durchführen, da sich auf der ersten Stufe bereits verschiedene Ausdrücke für s' und S' ergeben, diese aber nach Annahme J gleich sein müssen. Benutzt man dagegen Annahme H , so entsteht wieder (44).

Wir können jetzt das Resultat betrachten. Die Doppelgabel ist ihrer Verkettung nach symmetrisch für Vergangenheit und Zukunft, trotzdem ergibt sich ein charakteristischer Unterschied in der Art der Wahrscheinlichkeit. Die Vorwahrscheinlichkeit berechnet sich nach (40) zu $w_1 = p_1 q_1 P_1 Q_1$, für die Rückwahrscheinlichkeit aber wird dieser Ausdruck noch durch einen Summenausdruck dividiert. Und zwar gilt dies sowohl für Annahme J nach (39) als auch für Annahme H nach (44). Besonders deutlich wird diese Unsymmetrie für Annahme H , welche die Deutung der p'_i , q'_i , P_i , Q_i als Einzelrückwahrscheinlichkeiten zuläßt. Fassen wir nämlich die Doppelgabel in der umgekehrten Zeitrichtung auf, also C und D als frühere, A_1 und B_1 als spätere Ereignisse, die gestrichenen Einzelwahrscheinlichkeiten als Vorwahrscheinlichkeiten, die ungestrichenen als Rückwahrscheinlichkeiten, so würde sich

$$W(CD \rightarrow A_1 B_1) = w'_1 = p'_1 q'_1 P'_1 Q'_1 \quad (45)$$

$$W(A_1 B_1 \rightarrow CD) = w_1 = \frac{p_1 q_1 P_1 Q_1}{\sum_1^m \sum_1^n p_i q_k P_i Q_k} \quad (46)$$

ergeben.¹⁾ Das Maß der Wahrscheinlichkeitsimplikation zwischen den Ereignissen $A_1 B_1$ einerseits und CD andererseits wäre also ein anderes, obgleich die Wahrscheinlichkeiten zwischen den einzelnen Ereignissen ihren Zahlwert behalten hätten. Denken wir uns, es sei nicht bekannt, in welcher Zeitrichtung die Doppelgabel aufzufassen wäre, dagegen seien die Einzelwahrscheinlichkeiten in beiden Richtungen, also $p_1 q_1 P_1 Q_1$ und $p'_1 q'_1 P'_1 Q'_1$, und außerdem die Gesamtwahrscheinlichkeiten w_1 und w'_1 bekannt, so läßt sich die Zeitrichtung der Gabel bestimmen, indem man nachsieht, ob diese Größen die Beziehungen (40) und (44) oder (45) und (46) erfüllen. Für die vollständige Prüfung ist auch noch die Kenntnis der Wahrscheinlichkeiten mit von 1 verschiedenem Index erforderlich, die ja gerade so bestimmt werden können wie die anderen; aber auch ohne diese Größen läßt sich schon entscheiden, ob w_1 durch (40) oder w'_1 durch (45) richtig wiedergegeben wird. Die Richtung der Zeit läßt sich also durch Messung von Wahrscheinlichkeitsimplikationen, d. h. grundsätzlich durch Auszählung statistischer Regelmäßigkeiten, bestimmen.

Dagegen ist nichts darüber ausgesagt, ob der Vergangenheitsschluß die größere oder kleinere Wahrscheinlichkeit liefert. Nach (44), also Annahme H , ist

$$w'_1 \begin{matrix} \geq \\ < \end{matrix} w_1, \text{ wenn } \sum_1^m \sum_1^n p_i q_k P_i Q_k \begin{matrix} \leq \\ \geq \end{matrix} 1 \quad (47)$$

¹⁾ Hier beziehen sich die Wahrscheinlichkeiten mit von 1 verschiedenem Index auf Implikationen zwischen A_1 und B_1 einerseits und Ereignissen C_i und D_k andererseits, die wir vorher nicht betrachtet haben; diese Wahrscheinlichkeiten sind also mit solchen der vorangehenden Formeln nicht vergleichbar. Dagegen beziehen sich alle Wahrscheinlichkeiten mit dem Index 1 auf dieselben Implikationen wie vorher, und zwar in derselben Richtung zwischen den Ereignissen, nur daß, da wir die Ereignisse jetzt zeitlich umgekehrt annehmen, die gestrichenen Größen sich jetzt auf die Richtung „zeitlich vorwärts“ beziehen.

Bei Annahme J liefert (39)

$$w'_1 \underset{\leq}{\overset{\geq}{=}} w_1, \text{ wenn } \sum_1^m p_i q_i P_i Q_i \underset{\leq}{\overset{\geq}{=}} 1 \quad (48)$$

Über diese beiden Beziehungen aber ist allgemein nichts auszusagen, da die p_i , q_i , P_i , Q_i unverbundene Wahrscheinlichkeiten darstellen. Man erkennt auch hier wieder: nicht der Grad der Bestimmbarkeit, sondern die Art der Bestimmbarkeit unterscheidet Vergangenheit und Zukunft. Unter Umständen kann der Schluß in die Zukunft sicherer sein als der Schluß in die Vergangenheit, verglichen an derselben identischen Kausalverkettung.

Dagegen ist nach (44)

$$w'_1 > p'_i q'_i P_i Q_i \quad (49)$$

denn

$$\sum_1^n \sum_1^n p'_i q'_k P_i Q_k < \sum_1^m \sum_1^n \sum_1^m \sum_1^n p'_i q'_k P_r Q_s = \sum_1^m p'_i \sum_1^n q'_k \sum_1^m P_r \sum_1^n Q_s = 1$$

nach (41). Darum ist das durch (44) bestimmte w'_1 größer als das durch (45) bestimmte, und wir dürfen sagen: wenn man eine Kausalverkettung V_1 nach Fig. 15 mit einer andern V_2 vergleicht, die mit der ersten spiegelbildlich symmetrisch ist, gespiegelt an dem Gegenwartsquerschnitt, ergibt sich für den Vergangenheitschluß in V_1 ein höherer Grad der Wahrscheinlichkeit als für den Zukunftsschluß in V_2 . Nur in dieser Beziehung bedeutet der Unterschied in der Art des Schlusses auch einen Unterschied im Grad der Bestimmbarkeit.

Damit bestätigt sich auch für die Doppelgabel das Resultat, das wir schon für die Sattलगabel im Vergleich zur Spitzgabel gewonnen hatten. Für den Vergleich dieser Gabeln muß man noch einen Unterschied beachten. Der Vergangenheitschluß der Sattलगabel ist ein Schluß von zwei Ereignissen auf eines; er muß spiegelbildlich mit dem Zukunftsschluß (25) der Spitzgabel verglichen werden und ergibt dann in (33) und (37) das Resultat, das wir in (34) formuliert haben. Hier gilt also ebenfalls die größere Sicherheit für den Vergangenheitschluß. Der Vergangenheitschluß der Spitzgabel ist dagegen ein Schluß von einem Ereignis auf zwei; er muß spiegelbildlich mit dem Zukunftsschluß der Sattलगabel verglichen werden. Für den letzteren wird $r_1 = p_1 \cdot q_1$; vergleicht man dies mit (31) und (30), indem man in

letzteren Formeln die Bildungen aus den gestrichenen Größen der Bildung $p_1 \cdot q_1$ gegenüberstellt, so ergibt sich nur für Annahme J die größere Sicherheit des Vergangenheitsschlusses, während für Annahme H sich Gleichheit ergibt. Dieser Fall bildet also eine Ausnahme von unserer Regel.

Das Resultat des Abschnitts III dürfen wir in die Sätze zusammenfassen:

1. Die Existenz von Vor- und Rückwahrscheinlichkeit bedingt ein Verteilungsgesetz für Ereignisse in der Welt.

2. Das Maß der Rückwahrscheinlichkeit ist aus dem der Vorwahrscheinlichkeit nicht ohne zusätzliche Annahmen zu berechnen, die für die einzelnen Arten von Fällen verschieden sein können, und über deren Stattfinden empirisch entschieden werden muß.

3. Es läßt sich nicht sagen, daß die Vergangenheit in einer vorliegenden Weltstruktur stets mit größerer Wahrscheinlichkeit bestimmt werden kann als die Zukunft. Aber die topologische Besonderheit des Vergangenheitsschlusses bewirkt in den behandelten Fällen (abgesehen von einer Ausnahme), daß sich die Vergangenheit aus gegebenen Einzelrückwahrscheinlichkeiten sicherer berechnen läßt, als wenn dieselben Einzelwahrscheinlichkeiten Vorwahrscheinlichkeiten wären und zu einem spiegelbildlich symmetrischen Schluß in die Zukunft zu kombinieren wären.
