

2.1 Begriff und zahlenmäßiger Ausdruck des Einzelrisikos

Versicherer verpflichten sich im Versicherungsvertrag gegenüber dem Versicherungsnehmer zu vorab klar definierten Leistungen, die dem Zufall unterliegen. Somit ist die Fälligkeit der Leistungen im Voraus unsicher. Die Unsicherheit kann sich darauf beziehen,

- ob,
- wann oder
- in welcher Höhe

die Leistungen zu erbringen sind.

Im Regelfall erstreckt sich die Unsicherheit auf alle drei Aspekte. Dies muss jedoch nicht so sein. Ein Beispiel stellt die gemischte Lebensversicherung (Kapitallebensversicherung) dar. In ihrer Grundform ist die Höhe der Versicherungsleistung vertraglich fixiert. Es werden sowohl eine Erlebensfall- als auch eine Todesfallleistung vereinbart. Somit ist sicher, dass das Versicherungsunternehmen während der Vertragslaufzeit leisten muss. Die Fälligkeit ist somit nur noch bezüglich des Zeitpunktes unsicher.

Aus der Perspektive des Versicherers bezeichnen wir als **Einzelrisiko** die Ungewissheit der nach Eintritt und/oder Höhe zufälligen (vertraglichen) Leistungen. Dabei interessiert zunächst nur, welche Eigenschaften und Wirkungen den Versicherer betreffen und wie sie operational so zu erfassen und auszudrücken sind, dass sie in wirtschaftliches Handeln einbezogen werden können.

Die möglichen Leistungen lassen sich in Geldeinheiten messen, denn üblicherweise schulden Versicherer im Versicherungsfall dem Versicherungsnehmer eine Geldleistung. In einigen Sparten, wie z. B. der Kfz-, der Gebäude oder der Hausratversicherung, verpflichten sich Versicherer auch zu Naturalleistungen (z. B. Assistanceleistungen). Dies bedeutet, dass der Versicherer den Schaden beheben lässt, ohne dass der Geschädigte eine

direkte Geldleistung erhält. Eine solche Leistung kann unschwer durch einen gleichwertigen Geldbetrag gemessen werden.

Die vertraglich zugesagten Beträge stellen jedoch bedingte Leistungen dar, die an ein zufälliges Ereignis, den Eintritt eines **Versicherungsfalls** (Abschn. 5.2), gebunden sind. Wenn der Versicherungsfall eintritt, hat das Versicherungsunternehmen zu leisten.

Generell knüpft der im Vertrag definierte Versicherungsfall an ein Ereignis an, das den Versicherten „bedroht“; das heißt, das ihn in seiner wirtschaftlichen Situation oder allgemein beeinträchtigen kann (Schadenereignis). Die Zufallsereignisse sind für den Versicherer externer Natur, da ihr Eintritt auch ohne Bestehen eines Versicherungsvertrages möglich ist. Sie werden für den Versicherer erst durch den Versicherungsvertrag relevant. Der Versicherungsfall als das die Versicherungsleistung auslösende Ereignis wird erst durch den Vertrag konstituiert und gegenüber dem äußeren Ereignis abgegrenzt. Ein Gegenstück dazu wäre der Kauf bzw. Verkauf eines Lotterieloses. Der Loskäufer befindet sich im Gegensatz zu dem Versicherten vor dem Kauf nicht in einer Risikosituation. Die Risikosituation des Lotterieunternehmers ist zwar der des Versicherers ähnlich, aber er schafft erst das Zufallsexperiment, das die leistungsauslösenden Ereignisse hervorbringt, und kann daher mit festen Wahrscheinlichkeiten rechnen.

Wenn wir vom Schadenereignis gesprochen haben, so gilt dies in einem weiteren Sinne. Der Sinn der Versicherungsnachfrage liegt darin, die Risikosituation eines Versicherten durch kompensatorische Leistungen zu verbessern. Man sagt, Versicherung diene der Bedarfsdeckung. Werden, wie beim Pferderennen, Zahlungen bei Eintritt externer zufälliger Ereignisse versprochen, die den Empfänger nicht negativ berühren, so ist das als eine Wette einzustufen. Wettet man jedoch als Anhänger des FC Bayern München vor der Saison darauf, dass der FC Bayern München *nicht* Deutscher Meister wird, so wirkt die Wette wie eine Versicherung, da durch diese Wette eine negative Konsequenz beim Zahlungsempfänger abgesichert wird. Das Beispiel verdeutlicht, dass die Grenze zwischen einer Wette und einer Versicherung insbesondere vom sog. **versicherten Interesse** abhängt. Der Versicherungsnehmer versucht, mögliche negative Konsequenzen durch eine Versicherung abzumildern. Das versicherte Interesse ist somit eine Grundvoraussetzung für die Versicherungsnachfrage. In § 80 Abs. 1 und 2 VVG ist beispielsweise geregelt, dass ein Versicherungsnehmer dann nicht mehr zur Zahlung der Versicherungsprämie verpflichtet ist, wenn das versicherte Interesse im Verlauf des Vertragsverhältnisses entfällt. Natürlich endet in diesem Fall der nachträgliche Interessenmangel der Versicherungsvertrag. Dies kann etwa dann der Fall sein, wenn die versicherte Sache, z. B. ein Kfz, verkauft wurde und dafür keine Haftpflichtversicherung mehr benötigt wird. Grundsätzlich kann Versicherungsschutz bei Vorliegen eines Interessenmangels zu adversen Anreizen führen. Dies werden wir nachfolgend an zwei Beispielen (Beispiel 2.1 und Beispiel 2.2) illustrieren.

Beispiel 2.1

Black et al. (2015) verdeutlichen die Relevanz des versicherten Interesses an einem amerikanischen Beispiel aus den Fünfzigerjahren: Eine angeheiratete Tante hatte, als Versicherungsnehmerin und Begünstigte, mehrere Lebensversicherungen auf das Le-

ben eines Kindes abgeschlossen, das sie später tötete. Der Vater des Kindes verklagte mit Erfolg die Lebensversicherer, die es versäumt hatten, zu klären, ob ein versichertes Interesse bestand. Nach Ansicht des Gerichts wurde dadurch das Mordmotiv erzeugt.

Beispiel 2.2

Die Frage, ob ein versichertes Interesse vorliegt, hat nicht nur im Zusammenhang mit Versicherungsverträgen eine große Bedeutung, sondern analog auch für andere verwandte Finanztransaktionen. Dies hat z. B. im Zusammenhang mit der europäischen Staatsschuldenkrise eine hohe praktische Relevanz: Staatsanleihen einiger europäischer Staaten sind durch deren hohe Verschuldung vom Ausfall bedroht. Gegen das Risiko, dass der Emittent einer Anleihe seinen Zahlungsverpflichtungen nicht nachkommt, kann man sich mittlerweile am Kapitalmarkt mit Hilfe eines „Credit Default Swaps“ (CDS, „Kreditsausfall-Swap“), absichern. Dieses Kreditderivat kann dabei wie eine Versicherung wirken. Im Rahmen einer Transaktion zahlt der sog. Sicherungsnehmer eine laufend zu entrichtende bzw. eine einmalige Prämie und erhält dafür vom sog. Sicherungsgeber eine Ausgleichszahlung, falls die Anleihe des betreffenden Referenzschuldners nicht oder nur teilweise bedient wird. Der Wert eines solchen CDS steigt, wenn die Ausfallwahrscheinlichkeit der betreffenden Staatsanleihe steigt. Somit hat der Besitzer eines CDS Anreize, durch entsprechende (Falsch-)Informationen Markteinschätzungen bezüglich des Ausfallrisikos zu beeinflussen, wenn er die entsprechende Staatsanleihe nicht besitzt. In diesem Fall wirkt das Kreditderivat nicht wie eine Versicherung, sondern dient zu Spekulationszwecken.

Wir können festhalten, dass das Einzelrisiko vollständig durch die möglichen Zufallsergebnisse (Versicherungsfälle, welche aus externen Schadenereignissen abgeleitet wurden) und die zugehörigen Versicherungsleistungen festgelegt ist. Unabdingbar für die Versicherungsentscheidung – sowohl des Versicherers als auch des Versicherungsnehmers – ist eine Einschätzung der Wahrscheinlichkeiten der möglichen Ereignisse bzw. Versicherungsfälle. Sie gilt es, genau zu quantifizieren. Es kommt letztlich nicht darauf an, wie und warum Ereignisse auftreten, auch wenn nähere Kenntnisse für die praktische Gestaltung und Bewertung der Risiken von großer Bedeutung sind. Ein Gebäudeversicherer kann sich zum Beispiel aktuell auf das Wissen beschränken, mit welcher Wahrscheinlichkeit gravierende Stürme über bestimmte Regionen hinwegziehen und welche potenziellen Schäden verursacht werden. Nähere Informationen über relevante Einflussfaktoren, wie z. B. Klimawandel, machen es dem Versicherer aber möglicherweise langfristig einfacher, am Markt zu bestehen.

Letztlich wird jeder Versicherungsvertrag, mit den im Vertrag festgelegten Leistungen des Versicherers, durch eine Zufallsvariable repräsentiert. Das im Vertrag übernommene Einzelrisiko hat demnach zwei Dimensionen – die Leistungsbeträge und ihre Wahrscheinlichkeiten – und wird durch die Wahrscheinlichkeitsverteilung der Versicherungsleistungen gut beschrieben. Diese stellt das entscheidungsrelevante zahlenmäßige Abbild des Einzelrisikos dar.

Vor allem im Hinblick auf die Messung eines Einzelrisikos stellt sich die Frage, wie man geeignete Schadenwahrscheinlichkeiten bestimmen kann. Da Informationen über Einzelrisiken stets unvollständig sind, ist, wie wir in Abschn. 2.3.3 noch sehen werden, die Zuordnung von Wahrscheinlichkeiten zu Ereignissen letztlich immer subjektiv und beruht auf individuellen Einschätzungen. Diese Einschätzungen hängen natürlich vom Wissensstand und der Informationsverarbeitung der beteiligten Personen ab. Somit können unterschiedliche Personen demselben Ereignis auch unterschiedliche subjektive Wahrscheinlichkeiten zuordnen. Für Versicherungsunternehmen bedeutet dies, dass nicht zwingend umfassende Schadenstatistiken vorliegen müssen, um Wahrscheinlichkeiten von Schadenereignissen zu quantifizieren.

Durch die Quantifizierung und die Abbildung von Einzelrisiken als Zufallsvariablen bzw. Wahrscheinlichkeitsverteilungen haben wir versicherungstechnische Entscheidungen des Versicherers und Versicherungsnachfrageentscheidungen des potenziellen Versicherungsnehmers auf die Wahl zwischen unterschiedlichen Wahrscheinlichkeitsverteilungen reduziert. Dabei würden idealerweise immer alle Aspekte einer Verteilung Berücksichtigung finden. Allerdings zeigt sich, dass es nur unter ganz speziellen Voraussetzungen möglich ist, ganze Verteilungen miteinander zu vergleichen. Deshalb muss man sich in der Praxis in der Regel damit begnügen, bestimmte Parameter bzw. Kennzahlen einer Verteilung, wie den Erwartungswert und die Varianz (bzw. Standardabweichung), zu betrachten (Abschn. 3.2).

Wenn in der Praxis häufig das versicherte Objekt (z. B. das Fahrzeug in der Kfz-Versicherung) oder die versicherte Person (z. B. in der Unfallversicherung) als das Risiko bezeichnet werden, so steht das nicht im Widerspruch zu unserer Definition; denn gemeint ist die Zufallsvariable der damit verbundenen Schadenmöglichkeiten. Der Ausdruck Risiko soll das betreffende Objekt gerade und ausschließlich in seiner Eigenschaft als Quelle von unsicheren Schadenergebnissen charakterisieren.

2.2 Praktische Abgrenzung des Einzelrisikos

Wenn wir das Einzelrisiko als die kleinste Einheit für versicherungstechnische Entscheidungen betrachten, ist es nicht immer sinnvoll, unmittelbar den Versicherungsvertrag zugrunde zu legen. Die Abgrenzung der juristischen Einheit Vertrag bestimmt sich nach anderen als rein versicherungstechnischen Gesichtspunkten.

Beispiele dafür, dass es zweckmäßig sein kann, innerhalb eines Vertrages mehrere Einzelrisiken zu unterscheiden, sind die Krankenversicherung einer Familie oder die Feuerversicherung eines weitläufigen Industriebetriebes. Im ersten Fall ist es sinnvoll, die Einzelrisiken auf die Personen zu beziehen, im zweiten Fall auf die Feuerkomplexe. In einer verbundenen Versicherung (z. B. Hausrat) empfiehlt es sich, die Einzelrisiken nach den versicherten Gefahren Feuer, Einbruchdiebstahl etc. zu unterscheiden. Andererseits kann es auch vorteilhaft sein, mehrere Verträge für die Bildung eines Einzelrisikos zusam-

menzufassen, z. B. wenn auf einem Betriebsgrundstück Gebäude und Inhalt in getrennten Verträgen gegen Feuer versichert sind.

Welche Abgrenzung tatsächlich getroffen werden soll, lässt sich nicht allgemeingültig ableiten. Es kommt für die Zweckmäßigkeit der Abgrenzung auch auf die Art der zu treffenden Entscheidung an. Es kann sehr wohl sein, dass für die Prämienkalkulation andere Einheiten geeignet erscheinen als für die Rückversicherungspolitik. Allgemein lässt sich theoretisch nur empfehlen, die Einzelrisiken so abzugrenzen, dass sie einerseits aus möglichst einheitlichen, geschlossenen Ursachensystemen entstehen und in sich homogen erscheinen, dass sie aber andererseits isoliert und unabhängig von anderen Einzelrisiken sich gegenseitig möglichst nicht beeinflussen.

2.3 Die formale Struktur der Entscheidung über das Einzelrisiko

2.3.1 Das Grundmodell der Entscheidungstheorie

Zur Veranschaulichung der bedeutsamen Elemente des Einzelrisikos verwenden wir das **Grundmodell der Entscheidungstheorie** für die Versicherungsentscheidung. Mit dessen Hilfe kann das Entscheidungsproblem der für eine Versicherungsbeziehung relevanten Entscheider (Versicherungsunternehmen und Versicherungsnehmer) anschaulich dargestellt und einfach analysiert werden.

Für das Grundmodell der Entscheidungstheorie werden folgende Größen benötigt:

- 1) Ein **Aktionsraum** (A), der die Menge aller zur Verfügung stehenden Handlungsalternativen a_i mit $i = 1, \dots, n$ umfasst.
- 2) Ein **Zustandsraum** (Ω), bestehend aus allen vom Entscheider für möglich gehaltenen und für die Entscheidung relevanten Umweltzuständen ω_j mit $j = 1, \dots, m$.
- 3) Ein **Ergebnisraum** (E), die Menge aller für möglich erachteten Ergebnisse e_{ij} .
- 4) Eine **Ergebnisfunktion** $f: A \times \Omega \rightarrow E$, die jedem Paar (a_i, ω_j) ein Ergebnis e_{ij} zuordnet: $e_{ij} = f(a_i, \omega_j)$.

Zu den einzelnen Modellgrößen ist Folgendes anzumerken:

Zu 1) Es ist offensichtlich, dass ein Individuum nur Handlungsmöglichkeiten aufführen wird, deren Realisierung im Prinzip möglich ist. So werden Menschen beispielsweise bei der Wohnungssuche Handlungsalternativen nicht berücksichtigen, die außerhalb ihres von Restriktionen (wie z. B. der finanziellen Situation) festgelegten Rahmens liegen. Des Weiteren ist die Auswahl der Alternativen auch abhängig vom Informationsstand der Individuen. So wird ein Steuerberater zur Verminderung seiner Steuerlast Handlungsmöglichkeiten in Betracht ziehen, an die andere Individuen gar nicht denken würden. Unabhängig davon ist für Letztere die Handlungsmöglichkeit „Informationsbeschaffung durch

Heranziehen eines Steuerberaters“ relevant. Die Festlegung des Aktionsraumes ist somit nur subjektiv möglich. Weiterhin ist zu sagen, dass der Aktionsraum in einem Entscheidungsmodell mehrelementig sein muss, da sonst kein Entscheidungsproblem vorliegt. Außerdem müssen die einzelnen Handlungsalternativen sich gegenseitig ausschließen. Weitere Anforderungen an den Aktionsraum werden nicht gestellt.

Zu 2) Unter Umweltzuständen versteht man einander ausschließende Konstellationen von Ausprägungen der entscheidungsrelevanten Daten, die der Entscheider für möglich erachtet. Die Menge der Umweltzustände kann sowohl endlich als auch unendlich sein. Dabei liegen sowohl die Festlegung der entscheidungsrelevanten Daten selbst als auch ihre für möglich gehaltenen Ausprägungen im subjektiven Ermessen des Entscheidungsträgers. Es ist daher ohne weiteres möglich, dass unterschiedliche Entscheidungsträger bei ein und demselben Entscheidungsproblem eine unterschiedliche Festlegung der entscheidungsrelevanten Umweltzustände vornehmen.

Von vielen Autoren werden Entscheidungsmodelle unter Unsicherheit nach den in ihnen enthaltenen Wahrscheinlichkeitsprämissen untergliedert. Ist der Entscheider in der Lage, den von ihm als relevant erachteten Umweltzuständen Wahrscheinlichkeiten zuzuordnen, so wird von Entscheidungen unter Risiko gesprochen. Dagegen bezeichnen Entscheidungen unter Ungewissheit Situationen, in denen der Entscheidungsträger keinerlei Vorstellung besitzt, mit welcher Wahrscheinlichkeit die Umweltzustände eintreten werden.¹ Wir werden uns nachfolgend mit dieser Unterscheidung eingehender beschäftigen.

Zu 3) Die Menge E ist die Gesamtheit aller möglichen Handlungskonsequenzen. Nun steht zur Beschreibung von Handlungskonsequenzen auch bei einer Beschränkung auf objektiv messbare Sachverhalte immer noch eine solche Fülle unterschiedlicher Maßgrößen zur Auswahl, dass jede überschaubare Darstellung erwarteter Handlungsergebnisse eine Einschränkung auf bestimmte Merkmale erfordert. Welche diese Merkmale sind, hängt von den individuellen Zielvorstellungen des Entscheidenden ab. Damit ein Entscheider Handlungskonsequenzen anhand seiner Zielvorstellungen vergleichen kann, sollte die Merkmalsauswahl daher so bestimmt werden, dass die Ergebnisse alle relevanten Aspekte im Hinblick auf die vom Entscheidungsträger angestrebten Ziele berücksichtigen. Für unsere Fragestellung ist dies allerdings unproblematisch, da es bei Versicherungsverträgen letztlich immer um monetäre Konsequenzen geht. Zu beachten ist allerdings, dass die Ausgangsrisikosituation eines potenziellen Versicherungsnehmers durchaus maßgeblich durch nicht-monetäre Konsequenzen eines Schadens, wie z. B. die emotionalen Folgen des Verlusts eines Familienerbstückes, bestimmt werden kann, deren monetäres Äquiva-

¹ Die Terminologie ist im Hinblick auf Unsicherheit und Ungewissheit nicht immer einheitlich. Wir folgen hier der Zuordnung von Neus (2015, Abschn. 1.5). In der Versicherungsökonomie ist zunehmend auch der Begriff der „Ambiguität“ von Bedeutung, der oft, aber nicht immer, synonym zum Begriff der „Ungewissheit“ verwendet wird. Wir wollen uns hier mit dem angeführten Begriff der „Ungewissheit“ begnügen.

Tab. 2.1 Grundmodell der Entscheidungstheorie

	ω_1	...	ω_j	...	ω_m
a_1	e_{11}	...	e_{1j}	...	e_{1m}
...					
a_i	e_{i1}	...	e_{ij}	...	e_{im}
...					
a_n	e_{n1}	...	e_{nj}	...	e_{nm}

lent im Rahmen der Entscheidung über einen Versicherungsvertrag zu bestimmen ist. Es wird vorausgesetzt, dass der Entscheider eine vollständige und transitive Präferenzordnung² auf E besitzt. Diese Annahme ist allerdings offensichtlich unproblematisch, wenn alle Konsequenzen durch monetäre Größen beschrieben sind, da dann mehr Geld in der Regel weniger Geld vorgezogen wird.

Zu 4) Die Forderung, dass das Zusammentreffen einer Handlungsmöglichkeit und eines Umweltzustandes zu einem eindeutigen Ergebnis führt, kann durch eine geeignete Wahl des Zustandsraumes grundsätzlich immer erfüllt werden. Ist beispielsweise bei einer ersten Formulierung eines Entscheidungsmodells ein Ergebnis nicht eindeutig durch eine Handlungsalternative und einen Umweltzustand determiniert, so kann dieses Problem durch eine weitergehende Differenzierung des Zustandsraumes beseitigt werden.

Damit ist das Grundmodell der Entscheidungstheorie hinreichend charakterisiert. Tab. 2.1 liefert eine Darstellung in Form einer Entscheidungsmatrix für den Fall diskreter Mengen von Handlungsalternativen und Umweltzuständen.

2.3.2 Anforderungen an Wahrscheinlichkeitsmaße

In vielen Situationen sind die Ergebnisse wirtschaftlichen Handelns nicht deterministisch bzw. für einen Entscheidungsträger nicht als deterministisch erkennbar (Beispiele: Würfeln einer Zahl, Kursentwicklung einer Aktie oder Eintritt eines Schadens).

Die einzelnen möglichen zufälligen Ergebnisse werden üblicherweise als Elementarereignisse bezeichnet. Beim Würfeln sind die Elementarereignisse die Augenzahlen: „1“, „2“, ..., „6“. Der **Ereignisraum** oder auch **Zustandsraum** (Ω) ist die Vereinigung sämtlicher Elementarereignisse. Ω umfasst alle möglichen Ereignisse z , wobei ein Ereignis mehrere Elementarereignisse enthalten kann. Beim Würfel kann ein Ereignis beispiels-

² Im Fall einer vollständigen und transitiven Präferenzordnung kann der Entscheider in paarweisen Vergleichen zwischen Alternativen angeben, welche er präferiert, oder ob er zwischen den beiden indifferent ist. Darüber hinaus führen die paarweisen Vergleiche zu keinem Widerspruch (Anforderung der Transitivität), sodass aus den geäußerten Präferenzen eine widerspruchsfreie Ordnung über alle möglichen Ergebnisse abgeleitet werden kann (vgl. beispielsweise Laux et al. 2014, Abschn. 5.4.1).

weise eine bestimmte Augenzahl oder eine gerade bzw. ungerade Zahl sein. Ein Ereignis ist also immer eine Teilmenge von Ω . Diesen Ereignissen sollen nun Wahrscheinlichkeiten zugeordnet werden. Dies geschieht mittels der Axiomatisierung der Wahrscheinlichkeitstheorie durch Kolmogorov (1933):

Eine Funktion p , die jedem Ereignis $z \in \Omega$ eine reelle Zahl zuordnet, heißt **Wahrscheinlichkeitsmaß**, und $p(z)$ heißt Wahrscheinlichkeit des Ereignisses z , wenn gilt:

- **Axiom 1:** $0 \leq p(z) \leq 1$
- **Axiom 2:** $p(\Omega) = 1$
- **Axiom 3:** Für abzählbar viele disjunkte Ereignisse $z_1, z_2, \dots$ mit $z_i \cap z_j = \emptyset \forall i \neq j$ gilt: $p(z_1 \cup z_2 \cup \dots) = p(z_1) + p(z_2) + \dots = \sum_{i=1}^{\infty} p(z_i)$

Aus Axiom 1 folgt, dass Wahrscheinlichkeiten immer zwischen null und eins liegen müssen. Ist die Wahrscheinlichkeit eines Ereignisses null, kann dieses nicht eintreten. Ist die Wahrscheinlichkeit eines Ereignisses eins, tritt dieses mit Sicherheit ein. Aus Axiom 2 folgt, dass mit Sicherheit ein Ereignis aus dem Ereignisraum eintreten wird. Im Falle des Würfeln wird man sicher ein Ereignis aus dem Ereignisraum $\Omega = \{1, 2, \dots, 6\}$ würfeln. Axiom 3 besagt, dass sich die Wahrscheinlichkeit für das Eintreten einer Menge disjunkter Ereignisse einfach aus der Summe der Wahrscheinlichkeiten der einzelnen Ereignisse ergibt. Die Wahrscheinlichkeit, eine gerade Augenzahl zu würfeln, ist somit: $p(2 \cup 4 \cup 6) = p(2) + p(4) + p(6)$.

Im Folgenden werden einige Rechenregeln für Wahrscheinlichkeiten und grundlegende Definitionen aufgeführt. Es gilt:

- $p(\emptyset) = 0, p(\Omega) = 1$
- $p(z) + p(\neg z) = 1$ mit $\neg z = \Omega \setminus z$
- $p(z_1 \cup z_2) = p(z_1) + p(z_2) - p(z_1 \cap z_2)$

Die Regeln besagen, dass die Wahrscheinlichkeit dafür, dass kein Ereignis aus dem Ereignisraums eintritt, gleich null ist. Die Wahrscheinlichkeit dafür, dass ein Ereignis des Ereignisraums Ω eintritt, ist gleich 1. Betrachtet man die Wahrscheinlichkeiten für das Eintreten des Ereignisses z und dafür, dass z nicht eintritt, so ist die Summe beider Wahrscheinlichkeiten eins. $\neg z$ bezeichnet dabei „nicht z “, also den Ereignisraum ohne z . Allgemein ergibt sich die Wahrscheinlichkeit von zwei kombinierten Ereignissen als die Summe der Wahrscheinlichkeiten für die Ereignisse abzüglich der Wahrscheinlichkeit der Schnittmenge.

Mit Hilfe dieser Rechenregeln können wir nun einfach definieren, unter welchen Bedingungen zwei Ereignisse z_1 und z_2 **stochastisch unabhängig** sind. Dies gilt, wenn:

$$p(z_1 \cap z_2) = p(z_1) \cdot p(z_2).$$

Die **bedingte Wahrscheinlichkeit** $p(z_1|z_2)$ wird definiert als:

$$p(z_1|z_2) = \frac{p(z_1 \cap z_2)}{p(z_2)}.$$

Die bedingte Wahrscheinlichkeit gibt die Wahrscheinlichkeit des Eintretens eines Ereignisses z_1 an, unter der Bedingung, dass ein Ereignis z_2 eintritt. Für zwei stochastisch unabhängige Ereignisse z_1 und z_2 folgt unmittelbar:

$$p(z_1|z_2) = \frac{p(z_1 \cap z_2)}{p(z_2)} = \frac{p(z_1) \cdot p(z_2)}{p(z_2)} = p(z_1)$$

und entsprechend

$$p(z_2|z_1) = p(z_2).$$

Bei stochastisch unabhängigen Ereignissen ist die Wahrscheinlichkeit des Eintritts eines Ereignisses z_1 unabhängig vom Eintritt des Ereignisses z_2 . So ist z. B. die Wahrscheinlichkeit, beim Würfeln eine gerade Zahl zu werfen, unabhängig davon, ob im vorherigen Wurf eine gerade Zahl gefallen ist.

Die bedingten Wahrscheinlichkeiten führen uns zum **Satz von der totalen Wahrscheinlichkeit**.

Seien $B_i, i = 1, \dots, n$ paarweise disjunkte Ereignisse mit positiver Wahrscheinlichkeit, für die $\bigcup_{i=1}^n B_i = \Omega$ gilt. Für jedes Ereignis z gilt dann

$$p(z) = p(z|B_1) \cdot p(B_1) + \dots + p(z|B_n) \cdot p(B_n) = \sum_{i=1}^n p(z|B_i) \cdot p(B_i).$$

Mit Hilfe des Satzes von der totalen Wahrscheinlichkeit kann nun das für die Entscheidungstheorie überaus wichtige **Bayes-Theorem** formuliert werden.³

Seien $B_i, i = 1, \dots, n$ paarweise disjunkte Ereignisse mit positiver Wahrscheinlichkeit, für die $\bigcup_{i=1}^n B_i = \Omega$ gilt. Ist dann z ein Ereignis, so gilt für jedes $k \in \{1, \dots, n\}$:

$$p(B_k|z) = \frac{p(z|B_k) \cdot p(B_k)}{p(z)} = \frac{p(z|B_k) \cdot p(B_k)}{\sum_{i=1}^n p(z|B_i) \cdot p(B_i)}.$$

Dabei heißt $p(B_k)$ A-priori-Wahrscheinlichkeit von B_k und $p(B_k|z)$ A-posteriori-Wahrscheinlichkeit von B_k unter Kenntnis von z . Die Bedeutung des Bayes-Theorems liegt darin, dass mit seiner Hilfe eine Wahrscheinlichkeitsanpassung bei Vorliegen zusätzlicher Informationen erfolgen kann, wie Beispiel 2.3 verdeutlicht.

³ Nähere Erläuterungen, Beweise und Beispielaufgaben zu den beiden Sätzen finden sich beispielsweise in Behrend (2013, Kap. 4).

Beispiel 2.3

Angenommen, die Wahrscheinlichkeit, an einer Krankheit zu erkranken sei $p(B) = 0,003$. Mit einem Test kann nun eine Person feststellen, ob die Krankheit vorliegt. Die bedingte Wahrscheinlichkeit für einen positiven Test (Ereignis z) bei Vorliegen der Krankheit sei $p(z|B) = 0,98$. Bei einer nicht erkrankten Person ($\neg B$) führt der Test mit der Wahrscheinlichkeit $p(\neg z|\neg B) = 0,99$ zu einem negativen Ergebnis ($\neg z$). Wie groß ist nun die Wahrscheinlichkeit, erkrankt zu sein, wenn wir ein positives Testergebnis beobachten $p(B|z)$? Mit Hilfe der Bayes-Formel können wir diese berechnen:

$$\begin{aligned} p(B|z) &= \frac{p(z|B) \cdot p(B)}{p(z|B) \cdot p(B) + p(z|\neg B) \cdot p(\neg B)} = \frac{0,98 \cdot 0,003}{0,98 \cdot 0,003 + 0,01 \cdot 0,997} \\ &= \frac{0,00294}{0,01291} \approx 0,2277. \end{aligned}$$

Im Nenner steht die Wahrscheinlichkeit für ein positives Testergebnis. Diese beträgt ca. 1,3 %. Im Zähler steht nun die Wahrscheinlichkeit, dass man erkrankt und ein positives Testergebnis erhält, die bei ca. 0,3 % liegt. Die Wahrscheinlichkeit, tatsächlich erkrankt zu sein, gegeben ein positives Testergebnis liegt vor, beträgt somit ungefähr 23 % (in Anlehnung an Behrend 2013, S. 125)

Wahrscheinlichkeitseinschätzungen können also in den folgenden Schritten gebildet bzw. verfeinert werden:

1. Der Entscheider bildet gemäß seiner Einschätzung subjektive A-priori-Wahrscheinlichkeiten.
2. Bei Vorliegen neuer Informationen werden die A-priori-Wahrscheinlichkeiten modifiziert. In die subjektiven A-posteriori-Wahrscheinlichkeiten fließt dann neben den A-priori-Wahrscheinlichkeiten auch die „objektive“ Datenlage mit ein.

2.3.3 Wahrscheinlichkeitsinterpretationen

Formal ist der Begriff der Wahrscheinlichkeit durch die Kolmogorov-Axiome definiert. Eine ganz andere Frage ergibt sich im Zusammenhang mit der Interpretation und Messung von Wahrscheinlichkeiten. Im Wesentlichen finden sich in der Literatur drei verschiedene Wahrscheinlichkeitsinterpretationen:

1. logische (bzw. objektive A-priori-)Wahrscheinlichkeiten
2. frequentistische (bzw. objektive A-posteriori-)Wahrscheinlichkeiten
3. subjektive Wahrscheinlichkeiten

Risiko und Versicherungstechnik

Eine ökonomische Einführung

Karten, W.; Nell, M.; Richter, A.; Schiller, J.

2018, XI, 216 S. 36 Abb.,

ISBN: 978-3-658-06308-5