



universität
wien

DIPLOMARBEIT

Titel der Diplomarbeit

Behandlung von Schadensfällen -
Die mathematischen Prinzipien der
Schadenversicherung

angestrebter akademischer Grad

Magister der Naturwissenschaften (Mag.rer.nat)

Verfasser:	Manuel Feindert
Matrikelnummer:	0405572
Studienrichtung:	Lehramtsstudium UF Mathematik, UF Physik (A 190 406 412)
Betreuer:	ao. Univ.-Prof. Dr. Peter Raith

Wien, im Februar 2010

Danksagung

Mein besonderer Dank gilt Peter Raith, der mich nicht nur bei der vorliegenden Diplomarbeit hervorragend betreut hat, sondern bei dem ich auch während meiner gesamten Studienzeit mehrere Lehrveranstaltungen absolvieren konnte. Dadurch war es mir möglich, nicht nur Wissen in einzelnen mathematischen Teilgebieten zu erwerben, sondern auch größere Zusammenhänge zu verstehen. Außerdem konnte ich aus den Vorlesungen und Seminaren wertvolle Hinweise und Anregungen zur Behandlung verschiedener Themen im Unterricht mitnehmen.

Besonders bedanken möchte ich auch bei meinen Eltern, die mir die Absolvierung des Studiums ermöglicht haben und auf deren Unterstützung ich immer zählen konnte. Danken möchte ich auch meinen beiden Schwestern für die Motivation, die sie mir gegeben haben. Für die jahrelange Unterstützung bedanke ich mich auch herzlich bei meiner Großmutter.

Außerdem bedanke ich mich bei meinen StudienkollegInnen, mit denen ich während meines Studiums viele schöne Stunden verbringen konnte.

Inhaltsverzeichnis

1	Einleitung	7
2	Das Versicherungswesen und seine Grundlagen	11
2.1	Das Versicherungsverhältnis	11
2.2	Die Schadenversicherung und die Schadenversicherungsmathematik . . .	13
2.3	Das versicherungstechnische Risiko	14
2.3.1	Das Zufallsrisiko	14
2.3.2	Das Änderungsrisiko	15
2.3.3	Das Irrtumsrisiko	15
2.4	Der Ausgleich im Kollektiv	15
2.5	Wichtige mathematische Grundlagen	17
2.5.1	Die Gesetze der großen Zahlen	17
2.5.2	Die Normalverteilung und der zentrale Grenzwertsatz	18
3	Modellierung von Risiken	21
3.1	Individuelles Modell	22
3.1.1	Grundlegendes zum Modell	22
3.1.2	Daten zur Erstellung eines Modells	24
3.1.3	Homogene Bestände	25
3.1.4	Die Verteilung der Schadenshöhen und die Verteilung des Gesamt- schadens	30
3.1.5	Verteilungen zur Modellierung von Großschäden	50
3.1.6	Ein Modell für den Gesamtschaden mit der Logarithmischen Nor- malverteilung	52
3.1.7	Inhomogene Bestände	57
3.2	Kollektives Modell	60
3.2.1	Allgemeines zum Modell	60
3.2.2	Verteilungen der Schadenzahl	64

3.2.3	Die Verteilung der Schadenshöhen	68
3.2.4	Berechnung der Verteilung des Gesamtschadens	68
4	Die Berechnung von Prämien	83
4.1	Grundlegende Eigenschaften von Beständen	84
4.2	Prämienberechnung mit Prämienprinzipien	85
4.2.1	Das Nettoprämienprinzip	87
4.2.2	Aus dem Nettoprämienprinzip abgeleitete Prinzipien	89
4.2.3	Berechnung von Prämien mit Nutzenfunktionen	95
4.2.4	Das Esscherprinzip	102
4.2.5	Gewinnbeteiligung	105
4.2.6	Übersicht und Fazit	106
4.3	Erfahrungstarifizierung	107
4.3.1	Grundsätzliches zur Erfahrungstarifizierung	107
4.3.2	Erfahrungstarife anhand von Bonus-/Malus-Systemen	108
5	Maßnahmen zur Risikominderung	115
5.1	Rückversicherung	115
5.1.1	Proportionale Rückversicherung	117
5.1.2	Nichtproportionale Rückversicherung	121
5.1.3	Einsatzbereiche der verschiedenen Rückversicherungsarten	137
5.2	Selbstbehalte	137
6	Zusammenfassung und Ausblick	141
	Literaturverzeichnis	143
	Internetquellen	147
	Abbildungsverzeichnis	149
	Kurzzusammenfassung	151
	Abstract	153
	Lebenslauf	155

1 Einleitung

Das Versicherungswesen basiert auf der Idee, dass viele Personen Geld bezahlen, um im Falle eines bestimmten Ereignisses Zahlungen zu erhalten, mit denen entstandene Schäden bzw. Nachteile ausgeglichen werden können. Die Vorgehensweise, sich zu Gruppen zusammenzuschließen um Schäden und Nachteile einzelner Gruppenmitglieder auszugleichen, ist in der Menschheitsgeschichte schon früh aufgetreten. Zusammenschlüsse, bei denen Geld eingezahlt wurde, um in Schadensfällen finanzielle Entschädigungen zu erhalten, traten erstmals im 14. Jahrhundert in den oberitalienischen Seestädten auf, konkret ging es dort um den Schutz vor den Gefahren der Seefahrt und des Seehandels (vgl. [Koc88, S. 227]). Die Seeversicherung stellt damit den ersten Zusammenschluss zum Schadensausgleich auf einer kaufmännischen Grundlage dar und kann daher als älteste Form der Versicherung bezeichnet werden. Im Lauf der Zeit wurden Versicherungen für viele andere Bereiche geschaffen und das Versicherungswesen entwickelte sich zu einem wichtigen Wirtschaftszweig. So wurden in Österreich im Jahr 2008 rund 16,2 Milliarden Euro für Versicherungsschutz ausgegeben, alleine 7,3 Milliarden Euro davon bezogen sich auf den Bereich der Schaden- und Unfallversicherung. Umgerechnet auf die Einwohnerzahl bedeutet das, dass jederÖ ÖsterreicherIn im Jahr 2008 fast 2000 Euro an Versicherungsprämien bezahlt hat. Das Versicherungswesen leistet einen Beitrag von 5,77 % zum nominalen Bruttoinlandsprodukt und ist dadurch einer der bedeutendsten Wirtschaftszweige. (Die hier genannten Daten wurden entnommen aus: [2].) Neben der wirtschaftlichen Bedeutung des Versicherungswesens kann man ihm nach Peter Koch auch eine soziale Funktion (dadurch, dass auch hohe Schäden für einzelne Personen tragbar werden, weil sie von einem Versicherungsunternehmen übernommen werden) und eine ethische Funktion (z. B. indem nicht nur die VersicherungsnehmerInnen selbst geschützt werden, sondern auch andere Personen, wie beispielsweise in der KFZ-Haftpflichtversicherung) zuschreiben (vgl. [Koc05, S. 1]).

Einen wesentlichen Teil zum Funktionieren des Versicherungsprinzips trägt die Mathematik bei, denn Versicherungsfälle und die damit verbundenen Ausgaben des Versicherungsunternehmens treten zufällig auf, während die Einnahmen in Form der Prämien

bereits im Voraus festgelegt sind. Aber Versicherungsunternehmen sind keine Unternehmen, die Wetten abschließen und auf ihr Glück hoffen, sondern Unternehmen, die das Ziel verfolgen, zumindest so hohe Einnahmen wie Ausgaben zu haben. Um dies zu erreichen müssen mathematische Methoden angewendet werden, die das Risiko für das Versicherungsunternehmen berechenbar bzw. abschätzbar machen.

Diese Arbeit versucht, grundlegende mathematische Prinzipien der Schadenversicherung darzustellen, der Schwerpunkt wird dabei auf den Umgang des Versicherungsunternehmens mit den übernommenen Risiken gelegt, konkret werden die wahrscheinlichkeitstheoretischen Modelle der Schadenversicherung, die Bestimmung von Prämien und Möglichkeiten zur Minderung der übernommenen Risiken thematisiert.

Am Beginn der Arbeit wird das Versicherungswesen mit seinen Grundlagen behandelt, es werden dabei das Versicherungsverhältnis, sowie wichtige Grundprinzipien der Versicherung dargestellt. Da der Begriff der Schadenversicherung in der Literatur nicht eindeutig verwendet wird, wird festgelegt, was in dieser Arbeit darunter verstanden wird. Das versicherungstechnische Risiko, das Versicherungsunternehmen von anderen Unternehmen unterscheidet, und der Ausgleich im Kollektiv, eine Grundlage für das Funktionieren des Versicherungsprinzips, werden ebenfalls behandelt. Schließlich werden noch die wichtigsten mathematischen Grundlagen der Schadenversicherung genannt.

Dann werden die wahrscheinlichkeitstheoretischen Modelle der Schadenversicherung beschrieben, mit denen die Verteilung des jährlichen Gesamtschadens ermittelt werden soll. Grundsätzlich kann zwischen dem individuellen Modell, das von den jährlichen Gesamtschäden der einzelnen Versicherungsverträge ausgeht, und dem kollektiven Modell, das von den Schadenshöhen einzelner Schäden ausgeht, unterschieden werden. Für beide Modelltypen werden Wahrscheinlichkeitsverteilungen behandelt, die sich für die Beschreibung der Schadenshöhen eignen. Beim kollektiven Modell werden zusätzlich Verteilungen eingeführt, die zur Beschreibung der Schadenanzahl verwendet werden können. Zur Ermittlung der Gesamtschadenverteilung im kollektiven Modell werden verschiedene Möglichkeiten gezeigt, nämlich die explizite Berechnung dieser an einem ausgewählten Beispiel, die numerische Berechnung mit dem Panjer-Verfahren (sowohl für diskrete, als auch für kontinuierliche Schadenshöhenverteilungen) und die Approximation der Verteilung mit dem Monte-Carlo-Verfahren.

Anschließend wird auf die Berechnung von Prämien eingegangen. Nachdem grundsätzliche Eigenschaften von Versicherungsbeständen, nämlich die Stabilität in der Größe und die Stabilität in der Zeit, beschrieben wurden wird die Prämienberechnung mittels Prämienprinzipien behandelt: Bei dieser Art der Prämienkalkulation wird die Prämie an-

hand bestimmter Eigenschaften eines Versicherungsvertrages (oder eines ganzen Bestandes), wie beispielsweise dem erwarteten Schaden und dessen Streuung, festgelegt ohne dabei den individuellen Schadenverlauf des Vertrages zu beachten. Neben Prämienprinzipien, bei denen man die Prämie aufgrund wahrscheinlichkeitstheoretischer Eigenschaften erhält, wird auch auf Prämienprinzipien eingegangen, die den Nutzen für das Versicherungsunternehmen miteinbeziehen bzw. ökonomische Überlegungen beinhalten. Darüber hinaus wird die Erfahrungstarifizierung behandelt, die den individuellen Schadenverlauf eines Versicherungsvertrages zur Berechnung der Prämie heranzieht. Konkret wird dazu das Bonus-/Malus-System dargestellt, das beispielsweise in der österreichischen KFZ-Haftpflichtversicherung zur Anwendung kommt.

Nach der Behandlung versicherungsmathematischer Modelle und der Ermittlung von Prämien wird darauf eingegangen, wie das versicherungstechnische Risiko für das Versicherungsunternehmen gemindert werden kann. Dabei wird einerseits die Rückversicherung behandelt, bei der Versicherungsunternehmen selbst Versicherungsverträge abschließen, um sich gegen bestimmte Gefahren zu schützen, andererseits wird die Beteiligung von VersicherungsnehmerInnen mittels Selbstbehalten beschrieben.

2 Das Versicherungswesen und seine Grundlagen

2.1 Das Versicherungsverhältnis

Die Grundlage des Zustandekommens eines Versicherungsverhältnisses ist der Versicherungsvertrag, der zwischen einem Versicherungsunternehmen und einer Versicherungsnehmerin oder einem Versicherungsnehmer abgeschlossen wird. In [FW01] (diese Quelle stützt sich ihrerseits auf die Rechtsprechung des deutschen Bundesverwaltungsgerichts) wird der Versicherungsvertrag folgendermaßen charakterisiert:

»Ein Versicherungsvertrag liegt [...] vor, wenn gegen Entgelt für den Fall eines ungewissen Ereignisses bestimmte Leistungen übernommen werden, wobei das übernommene Risiko auf eine Vielzahl von Personen, die durch die gleiche Gefahr bedroht sind, verteilt ist und der Risikoübernahme eine Beitragskalkulation zu Grunde liegt, die auf dem Gesetz der großen Zahlen beruht.« [FW01, S. 713]

Ähnlich wird der Versicherungsvertrag von Thomas Mack bestimmt:

»Im Versicherungsvertrag (Police) verpflichtet sich das Versicherungsunternehmen gegen Erhalt eines vereinbarten, im Voraus fälligen Geldbetrags (Prämie), bei Eintritt von im Vertrag näher definierten ungewissen Ereignissen (Schäden) bestimmte, in ihrer Höhe meist vom betreffenden Ereignis abhängende Zahlungen an den Vertragspartner (Versicherungsnehmer) zu leisten, die den aus dem Ereignis resultierenden wirtschaftlichen Nachteil des Versicherungsnehmers reduzieren oder ausgleichen sollen.« [Mac02, S. 23]

Anhand der beiden gegebenen Definitionen eines Versicherungsvertrages lassen sich die Grundzüge und Prinzipien des Versicherungswesens angeben: Versicherbar ist ein Ereignis nur dann, wenn der Eintritt dieses Ereignisses ungewiss ist. In mathematischer

Ausdrucksweise heißt das, dass eine Versicherung nicht abgeschlossen werden kann, wenn ein Ereignis fast sicher eintritt bzw. fast sicher nicht eintritt, die Wahrscheinlichkeit für den Eintritt eines versicherbaren Ereignisses muss also im Intervall $(0, 1)$ liegen. Die geforderte Ungewissheit muss sich dabei nicht immer darauf beziehen, ob das Ereignis eintritt, sondern kann sich auch auf den Zeitpunkt des Eintretens beziehen, so ist beispielsweise in der Ablebensversicherung klar, dass eine versicherte Person irgendwann sterben wird, allerdings ist der Zeitpunkt des Todes ungewiss.

Eine weitere wichtige Grundlage des Funktionierens von Versicherungsschutz ist es, dass viele Personen, die von einer bestimmten Gefahr bedroht sind, einen Versicherungsvertrag abschließen. Würde ein Versicherungsunternehmen nur mit einer einzigen Person einen Vertrag abschließen, der dieser Person finanzielle Leistungen beim Eintritt eines bestimmten Ereignisses garantiert, dann könnte dieser Vertrag nicht als Versicherungsvertrag betrachtet werden, sondern würde eher ein Glücksspiel oder eine Wette darstellen. Bei der Übernahme von Risiken durch Versicherungsunternehmen spielt auch die (zumindest annähernde) Unabhängigkeit der Risiken eine Rolle: So kann kein Versicherungsvertrag abgeschlossen werden, wenn eine Gefahr zwar für viele Personen besteht, wenn diese im Falle des Eintritts der Gefahr aber alle davon betroffen wären. Versicherungsunternehmen müssen deshalb beim Abschließen von neuen Versicherungsverträgen darauf achten, dass die im Vertrag definierten Ereignisse (weitgehend) unabhängig von den Ereignissen der bisher abgeschlossenen Verträge sind. Praktisch kann das beispielsweise bedeuten, dass Gebäude, die gegen Hagel- oder Hochwasserschäden versichert werden, über einen größeren geographischen Bereich verteilt sein müssen. Allerdings ist in der Praxis eine vollständige Unabhängigkeit nur selten möglich und es müssen eventuell Kompromisse eingegangen werden, dennoch muss beachtet werden, dass die Unabhängigkeit einzelner Risiken eine Grundlage der Versicherung darstellt (vgl. [Far06, S. 39]).

Um Versicherungsschutz zu erhalten müssen VersicherungsnehmerInnen eine Prämie bezahlen, die beim Abschluss des Versicherungsvertrages bereits festgelegt ist. Für VersicherungsnehmerInnen heißt das, dass sie im Voraus wissen, wie viel sie bezahlen werden und das unabhängig davon, ob das im Vertrag genannte Ereignis eintritt oder nicht. Das Versicherungsunternehmen hingegen weiß nicht, ob es in Zukunft Zahlungen an VersicherungsnehmerInnen zu leisten hat bzw. wie hoch eventuelle Zahlungen sein werden. In [10, S. 1] wird Versicherung deswegen auch als die Übernahme ungewisser Zahlungen gegen eine feste Prämie bestimmt. Für die Dienstleistung des Versicherungsschutzes muss trotz der Ungewissheit, die für das Versicherungsunternehmen besteht, beim Abschluss

des Vertrages ein Preis festgelegt werden. Dieser Preis, den VersicherungsnehmerInnen dann in Form der Prämie entrichten, kann jedoch nicht wie der Preis eines Produktes in anderen Wirtschaftsbereichen bestimmt werden, da der Aufwand für das Versicherungsunternehmen bei Vertragsabschluss nicht klar ist, denn sowohl die Anzahl der zu leistenden Zahlungen als auch deren Höhe hängt vom Zufall ab. Damit das Versicherungsunternehmen dennoch einen Preis kalkulieren kann müssen Modelle erstellt werden, die Schätzungen der zu tätigen Zahlungen erlauben. Mit Hilfe dieser Modelle lassen sich dann die Prämien für einzelne VersicherungsnehmerInnen festlegen. In der ersten angegebenen Definition des Versicherungsvertrages wird als Grundlage zur Kalkulation von Prämien explizit das Gesetz der großen Zahlen genannt, der Einfluss dieses Gesetzes auf die Versicherungsmathematik wird in weiterer Folge noch genauer betrachtet.

2.2 Die Schadenversicherung und die Schadenversicherungsmathematik

In der Literatur werden Versicherungen auf viele Arten voneinander unterschieden, eine Möglichkeit ist die Unterscheidung nach der Versicherungsform in Summenversicherung und Schadenversicherung. Die Summenversicherung wird dabei als Versicherungsform verstanden, die die Deckung eines im Voraus abgeschätzten Bedarfs zum Ziel hat (abstrakte Bedarfsdeckung), während die Schadenversicherung den tatsächlichen Bedarf (konkrete Bedarfsdeckung) decken soll (vgl. [Sch05, S. 78]). Neben der Unterteilung in Schadenversicherung und Summenversicherung tritt in der Literatur aber auch eine Unterscheidung in Schadenversicherung und Personenversicherung auf. Im österreichischen Versicherungsvertragsgesetz (VersVG) hingegen wird unterschieden in Schadenversicherung, Krankenversicherung, Lebensversicherung und Unfallversicherung (vgl. [3]).

Der Begriff der Schadenversicherung wird also nicht einheitlich verwendet, denn beispielsweise kann die Krankenversicherung der Personenversicherung zugeordnet werden, sie stellt aber keine Summenversicherung dar, nach dem österreichischen Versicherungsvertragsgesetz ist sie aber dennoch keine Schadenversicherung.

Der dieser Arbeit zugrunde gelegte Begriff der Schadenversicherung soll jener des österreichischen Versicherungsvertragsgesetzes sein, das heißt, die Krankenversicherung und die Unfallversicherung werden hier nicht als Arten der Schadenversicherung verstanden. Das Versicherungsvertragsgesetz unterteilt die Schadenversicherung weiter in verschiedene Versicherungsarten, nämlich die Feuerversicherung, die Hagelversicherung, die Tierversicherung, die Transportversicherung, die Haftpflichtversicherung und die Rechts-

schutzversicherung (vgl. [3]).

Die Schadenversicherungsmathematik kann beschrieben werden als »*die Gesamtheit der Modelle und Methoden zur quantitativen Beschreibung von (Nichtlebens-)Versicherungen, bei denen der Zeitpunkt des Eintretens eines Versicherungsfalls und die Höhe der Versicherungsleistung zufällig sind*« [Hei88, S. 753]. Dazu verwendet die Schadenversicherungsmathematik vor allem Erkenntnisse aus der Wahrscheinlichkeitstheorie und der Statistik. In der Praxis sind aber auch andere mathematische Teilgebiete von Bedeutung, beispielsweise dann wenn es um effiziente numerische Berechnungen geht.

Konkret werden in dieser Arbeit das Erstellen von Modellen, die Berechnung von Prämien und Maßnahmen zur Risikominderung (Rückversicherung und Selbstbehalte) behandelt.

2.3 Das versicherungstechnische Risiko

Versicherungsunternehmen unterscheiden sich von sonstigen Wirtschaftsunternehmen dadurch, dass sie durch einen starken Einfluss des Zufalls nicht nur von den üblichen unternehmerischen Risiken betroffen sind, sondern auch vom sogenannten versicherungstechnischen Risiko. Nach Peter Albrecht und Edmund Schwake kann das versicherungstechnische Risiko folgendermaßen definiert werden:

»*Das versicherungstechnische Risiko ist die Gefahr, dass für einen bestimmten Zeitraum der Gesamtschaden des versicherten Bestandes die Summe der für die reine Risikoübernahme zur Verfügung stehenden Gesamtprämie und des vorhandenen Sicherheitskapitals übersteigt.*« [AS88, S. 652]

Die Gründe, dass der Gesamtschaden höher ist als die Summe der eingenommenen Prämien und des eventuell vorhandenen Sicherheitskapitals, können vielfältig sein, in [Far06, S. 83f] wird unterschieden in Zufallsrisiko, Änderungsrisiko und Irrtumsrisiko.

2.3.1 Das Zufallsrisiko

Unter dem Zufallsrisiko versteht man die zufällige Abweichung des Gesamtschadens von einer Schätzung des Gesamtschadens, die Abweichung kann dabei sowohl auf zufällige Abweichungen der Schadenshöhen, als auch auf zufällige Abweichungen der Schadenanzahl zurückgeführt werden (vgl. [Far06, S. 83] und [Sch05, S. 56]). Versicherungsunternehmen sind dem Zufallsrisiko ständig ausgeliefert, so ist es wahrscheinlichkeitstheoretisch beispielsweise durchaus möglich, dass in einer Versicherungsperiode sehr viele hohe

Schäden auftreten, während in der darauffolgenden Periode nur wenige, kleine Schäden auftreten, obwohl der Gesamtschaden in beiden Perioden als gleich hoch eingeschätzt wurde. Für die Praxis der Schadenversicherung besonders bedeutend im Zusammenhang mit dem Zufallsrisiko ist das Risiko von ungewöhnlich großen Schäden und das Risiko, dass besonders viele Versicherungsverträge von einem Ereignis betroffen sind (z. B. bei Hochwasser oder Hagel).

Obwohl das Zufallsrisiko nicht eliminiert werden kann, können Versicherungsunternehmen das Zufallsrisiko senken, indem sie viele unabhängige Verträge zusammenfassen und so einen „Ausgleich im Kollektiv“ bewirken, der im nächsten Abschnitt noch genauer besprochen wird.

2.3.2 Das Änderungsrisiko

Der Gesamtschaden eines Versicherungsunternehmens wird mit einer Wahrscheinlichkeitsverteilung beschrieben. In die Ermittlung dieser Verteilung fließen Daten ein, die sich auf die Vergangenheit beziehen und die mitunter in der Zukunft nicht mehr zutreffend sein werden. Das damit verbundene Risiko wird als Änderungsrisiko bezeichnet, Beispiele dafür sind Veränderungen im Bereich der Technik, wirtschaftliche Veränderungen oder eine Änderung von meteorologischen Gegebenheiten (vgl. [Sch05, S. 58]).

2.3.3 Das Irrtumsrisiko

Bei der Ermittlung der Verteilung des Gesamtschadens können Fehler passieren. Fehler, die auf falsche oder unvollständige Daten zurückgeführt werden können und Fehler, die durch einen fehlerhaften Umgang mit vorhandenen Daten entstehen, werden als Irrtumsrisiko bezeichnet (vgl. [Sch05, S. 58]). Beispiele sind die Verwendung veralteter Statistiken, eine falsche Interpretation von Daten oder die Berücksichtigung zu weniger Daten.

2.4 Der Ausgleich im Kollektiv

Eine Grundbedingung zum Funktionieren des Versicherungswesens ist der Ausgleich im Kollektiv. Darunter wird der Umstand verstanden, dass Versicherungsverträge, die schadenfrei bleiben oder nur kleine Schäden verursachen große Schäden anderer Verträge ausgleichen (vgl. [Far06, S. 47f]). Dieser Ausgleich kann aber nur dann funktionieren, wenn viele unabhängige Verträge zusammengefasst werden.

In der Literatur wird der Ausgleich im Kollektiv mathematisch unterschiedlich begründet, hier soll die Begründung mit dem Variationskoeffizienten erfolgen, der die Streuung einer Zufallsvariable relativ zum Erwartungswert angibt (diese Vorgehensweise wurde entnommen aus [Mac02, S. 24]): Dazu betrachtet man n identische verteilte Zufallsvariablen $X_1, X_2, \dots, X_n$ die einen endlichen Erwartungswert und eine endliche Varianz aufweisen, und jeweils den jährlichen Gesamtschaden eines Versicherungsvertrages angeben. Für den Gesamtschaden S dieser n Verträge gilt dann:

$$S = \sum_{i=1}^n X_i.$$

Da die Zufallsvariablen $X_1, X_2, \dots, X_n$ identisch verteilt und unabhängig sind lassen sich der Erwartungswert und die Varianz des Gesamtschadens durch den Erwartungswert und die Varianz einer einzelnen Zufallsvariable angeben:

$$\begin{aligned} E(S) &= n \cdot E(X_1), \\ V(S) &= n \cdot V(X_1). \end{aligned}$$

Für den Variationskoeffizienten von S gilt dann:

$$\text{VK}(S) = \frac{\sqrt{V(S)}}{E(S)} = \frac{\sqrt{n \cdot V(X_1)}}{n \cdot E(X_1)} = \frac{\sqrt{V(X_1)}}{\sqrt{n} \cdot E(X_1)}.$$

Für den Variationskoeffizienten der Zufallsvariable X_1 hingegen gilt:

$$\text{VK}(X_1) = \frac{\sqrt{V(X_1)}}{E(X_1)}.$$

Der Variationskoeffizient von S ist indirekt proportional zu $\sqrt{n}$, durch das Zusammenfassen vieler Verträge wird der Variationskoeffizient des Gesamtschadens also immer kleiner. (Hier muss beachtet werden, dass nur der Variationskoeffizient kleiner wird, absolut wird die Streuung beim Zusammenfassen vieler Verträge größer.)

Das Bilden von Beständen, also das Zusammenfassen mehrerer Versicherungsverträge, wird im Kapitel über die Modellierung von Risiken noch näher erläutert.

2.5 Wichtige mathematische Grundlagen

Nachdem die Grundprinzipien des Versicherungswesens behandelt wurden sollen nun die wichtigsten mathematischen Grundlagen genannt werden. Das sind konkret die Gesetze der großen Zahlen und der zentrale Grenzwertsatz.

2.5.1 Die Gesetze der großen Zahlen

Die Gesetze der großen Zahlen liefern Konvergenzaussagen für Folgen von Zufallsvariablen. Wegen ihrer Bedeutung für das Versicherungswesen werden sie auch als „Produktionsgesetze der Versicherungstechnik“ bezeichnet (vgl. z. B. [Mac02, S. 26]).

Bemerkung. Falls in dieser Arbeit Mengen, Summen oder Folgen von Zufallsvariablen betrachtet werden, dann wird stets angenommen, dass alle Zufallsvariablen auf dem gleichen Wahrscheinlichkeitsraum definiert wurden.

Zuerst soll das schwache Gesetz der großen Zahlen genannt werden:

Satz 2.1. (*schwaches Gesetz der großen Zahlen*)

Sei $X_1, X_2, \dots$ eine Folge von quadratisch integrierbaren, identisch verteilten, unabhängigen Zufallsvariablen mit dem Erwartungswert $E(X_i) =: \mu$ für $i = 1, 2, \dots$. Dann gilt für alle $\epsilon > 0$:

$$\lim_{n \rightarrow \infty} P \left(\left| \frac{1}{n} \sum_{i=1}^n X_i - \mu \right| \geq \epsilon \right) = 0.$$

Das schwache Gesetz der großen Zahlen besagt, dass bei einer Folge von quadratisch integrierbaren, identisch verteilten, unabhängigen Zufallsvariablen mit dem Erwartungswert μ das arithmetischen Mittel stochastisch gegen den Erwartungswert μ konvergiert. In der Literatur wird das schwache Gesetz der großen Zahlen auch gelegentlich als Begründung für den Ausgleich im Kollektiv verwendet, in dieser Arbeit wurde der Ausgleich im Kollektiv aber durch den Variationskoeffizienten des Gesamtschadens begründet, für den im Gegensatz zum schwachen Gesetz der großen Zahlen kein Grenzwert gebildet werden muss um einen Ausgleichseffekt zu erkennen. Betrachtet man die Zufallsvariablen $X_1, X_2, \dots$ als Schadenshöhen einzelner Versicherungsverträge, dann zeigt auch das schwache Gesetz der großen Zahlen, dass durch das Zusammenfassen vieler Verträge die Wahrscheinlichkeit, dass der durchschnittliche Schaden und der erwartete Schaden pro Vertrag stark voneinander abweichen, verringert wird.

Satz 2.2. *(starkes Gesetz der großen Zahlen)*

Sei $X_1, X_2, \dots$ eine Folge von integrierbaren, identisch verteilten, unabhängigen Zufallsvariablen mit dem Erwartungswert $E(X_i) =: \mu$ ($i = 1, 2, \dots$). Dann gilt:

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n X_i = \mu \quad \text{mit } P\text{-fast sicher.}$$

Das starke Gesetz der großen Zahlen besagt, dass bei einer Folge von integrierbaren, identisch verteilten, unabhängigen Zufallsvariablen mit dem Erwartungswert μ das arithmetische Mittel fast sicher gegen den Erwartungswert μ konvergiert. Nach [Mac02, S. 26] und [10, S. 10] kann diese Aussage sowohl für die Berechnung von Versicherungsprämien, als auch für die Schätzung des erwarteten Schadens verwendet werden: Betrachtet man die Zufallsvariablen $X_1, X_2, \dots$ als die Schadenshöhen eines Versicherungsvertrages in einzelnen Perioden, dann kann das Versicherungsunternehmen μ als „fairen Preis“ für die Risikoübernahme betrachten, da die durchschnittliche Zahlung pro Periode fast sicher gegen μ konvergiert (dies wird im Kapitel über die Prämienberechnung noch genauer diskutiert). Für die Schätzung des erwarteten Schadens einer Zufallsvariable besagt das starke Gesetz der großen Zahlen, dass das arithmetische Mittel unabhängiger Zufallsvariablen, die wie die zu schätzende Zufallsvariable verteilt sind, dazu benutzt werden kann.

2.5.2 Die Normalverteilung und der zentrale Grenzwertsatz

Die Normalverteilung kann zwar nicht zur Modellierung von Schadenshöhen verwendet werden (aber die mit der Normalverteilung eng verwandte Logarithmische Normalverteilung), dennoch ist sie für die Versicherungsmathematik von Bedeutung. Zusammen mit dem zentralen Grenzwertsatz wird die Normalverteilung in weiterer Folge unter anderem benötigt, um zu zeigen, dass das Versicherungsunternehmen die Prämie höher als den erwarteten Schaden ansetzen muss, um nicht mit einer hohen Wahrscheinlichkeit mehr für die Begleichung von Schäden auszugeben als durch die Prämien eingenommen wurde.

Die Normalverteilung kann folgendermaßen definiert werden:

Definition 2.1.

Als Normalverteilung $N(\mu, \sigma^2)$ mit den Parametern $\mu \in \mathbb{R}$ und $\sigma^2 \in \mathbb{R}, \sigma^2 > 0$ wird die

Wahrscheinlichkeitsverteilung mit der Dichte

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma^2} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

bezeichnet.

Die Verteilungsfunktion der Normalverteilung kann angegeben werden als:

$$F(x) = \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^x e^{-\frac{(t-\mu)^2}{2\sigma^2}} dt.$$

Für den Erwartungswert und die Varianz der Normalverteilung gilt:

$$E(X) = \mu,$$

$$V(X) = \sigma^2.$$

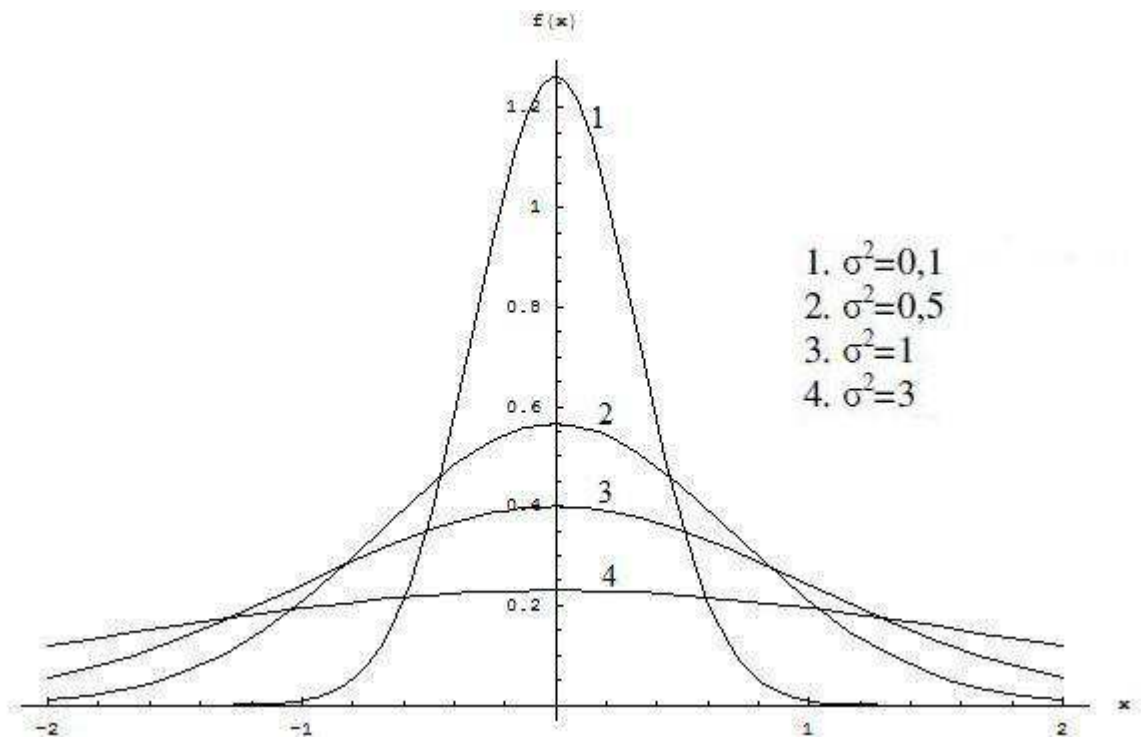


Abbildung 2.1: Die Dichtefunktion der Normalverteilung mit $\mu = 0$ und verschiedenen Werten für σ^2 .

Da die Verteilungsfunktion der Normalverteilung nicht integralfrei angegeben werden kann, werden in der Praxis die Funktionswerte der Verteilungsfunktion in Tabellen angegeben. Allerdings findet man dort üblicherweise nur die Werte für die Standardnormalverteilung (auch als 0-1-Normalverteilung bezeichnet) mit den Parameterwerten $\mu = 0$ und $\sigma^2 = 1$. Die Verteilungsfunktion der Standardnormalverteilung (die auch gaußsches Fehlerintegral genannt wird) wird mit Φ bezeichnet und kann folgendermaßen angegeben werden:

$$\Phi(x) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^x e^{-\frac{t^2}{2}} dt.$$

Für die Umrechnung der Verteilungsfunktionen der Standardnormalverteilung und der Normalverteilung mit beliebigen Werten für μ und σ^2 gilt:

$$F(x) = \Phi\left(\frac{x - \mu}{\sigma}\right).$$

Die große praktische Bedeutung der Normalverteilung beruht unter anderem auf dem zentralen Grenzwertsatz:

Satz 2.3. (*zentraler Grenzwertsatz*)

Sei $X_1, X_2, \dots$ eine Folge von quadratisch integrierbaren, identisch verteilten, unabhängigen Zufallsvariablen mit dem Erwartungswert $E(X_i) =: \mu$ und der Varianz $V(X_i) =: \sigma^2$ für $i = 1, 2, \dots$. Weiters sei die Zufallsvariable Z_n für $n = 1, 2, \dots$ folgendermaßen definiert:

$$Z_n := \frac{\sum_{i=1}^n X_i - n\mu}{\sqrt{n\sigma^2}}.$$

Dann gilt:

$$\lim_{n \rightarrow \infty} P(Z_n \leq z) = \Phi(z) \quad \forall z \in \mathbb{R}.$$

Im zentralen Grenzwertsatz wird eine Folge von quadratisch integrierbaren, identisch verteilten, unabhängigen Zufallsvariablen betrachtet. Der Satz besagt, dass die Verteilungsfunktion der Zufallsvariable Z_n , die sukzessive aus den einzelnen Folgengliedern gebildet wird, gegen die Verteilungsfunktion der Standardnormalverteilung konvergiert. Die Normalverteilung tritt deshalb in praktischen Anwendungen oft als Approximation anderer Verteilungen auf.

3 Modellierung von Risiken

Wie in zahlreichen anderen mathematischen Anwendungsgebieten geht es auch in der Schadenversicherungsmathematik darum, praktische Probleme in mathematischen Modellen zu beschreiben. Bei der Modellbildung kann man folgende Schritte unterscheiden (nach [Mac02, S.22]):

1. Wahl eines Modells,
2. Schätzung der Modellparameter,
3. Überprüfung der Übereinstimmung des Modells mit den Daten (Anpassungstest),
4. Berechnung der gesuchten Antwort auf die untersuchte Frage,
5. Validierung der Antwort, Prüfung auf Plausibilität und Sensitivität.

Die jeweils nachfolgenden Schritte setzen dabei den unmittelbar vorhergehenden Schritt voraus. Der wichtigste Schritt ist somit der erste, da auf ihm alle nachfolgenden Schritte basieren. Die Wahl eines nicht geeigneten Modells lässt sich in keinem der späteren Schritte korrigieren und ein nicht geeignetes Modell wird auch keine zufriedenstellenden Antworten für die modellierten Probleme liefern. Außerdem kann sich in jedem der Schritte herausstellen, dass es erforderlich ist, das Modell zu ändern. Dann ist es notwendig, wieder zum ersten Schritt zurückzugehen und alle folgenden Schritte müssen noch einmal bearbeitet werden. Gleichzeitig ist der erste Schritt oft aber auch der schwierigste, da es für die Modellwahl im Gegensatz zu den anderen Schritten nicht unbedingt immer geeignete Techniken gibt.

Es ist deshalb auf jeden Fall lohnend, genug Zeit für die Wahl des Modells aufzuwenden, da nur ein geeignetes Modell die untersuchten Fragen ausreichend beantworten kann. Thomas Mack charakterisiert ein geeignetes Modell folgendermaßen:

»Ein Modell ist geeignet, wenn es die für die Fragestellung wesentlichen Aspekte der Realität möglichst genau abbildet und dabei noch so einfach

bleibt, dass die Fragestellung mit erträglichem Aufwand innerhalb des Modells beantwortet werden kann. « [Mac02, S.22]

Bei der Modellierung von Risiken kann man folgende Modelltypen unterscheiden:

- Individuelles Modell,
- Kollektives Modell.

Beide Modelltypen werden im weiteren Verlauf beschrieben und diskutiert.

3.1 Individuelles Modell

3.1.1 Grundlegendes zum Modell

Der Ausgangspunkt des individuellen Modells sind einzelne Versicherungsverträge. Es werden die Schäden der einzelnen Verträge betrachtet, ausgehend von ihnen kommt man mittels Faltungsoperationen zu einer Verteilung des Gesamtschadens. Da ausnahmslos alle Verträge betrachtet werden müssen treten sehr viele Faltungsoperationen auf, was die Berechnung der Gesamtschadensverteilung aufwendig macht (vgl. [Gru96, S. 105]).

Versicherungsunternehmen übernehmen Risiken von VersicherungsnehmerInnen. Mathematisch kann man ein Risiko folgendermaßen beschreiben (vgl. [13, S. 4]):

Definition 3.1.

Ein Risiko ist eine Zufallsvariable X , für die gilt: $0 \leq X < \infty$.

Die Menge $\{X_1, X_2, \dots, X_n\}$ von Risiken wird Bestand genannt, die Zahl n heißt Größe des Bestandes.

Im individuellen Modell soll der Gesamtschaden eines Jahres aus einzelnen Risiken modelliert werden. Dazu werden mehrere Risiken zu einem Bestand (gelegentlich auch als Portfolio oder Kollektiv bezeichnet) zusammengefasst. Ein Risiko kann hier als unteilbare Einheit angesehen werden, dies kann eine Versicherungspolice oder beispielsweise ein Haus in einer Feuerversicherung sein (vgl. [Wol88, S. 50]). Die Zufallsvariable X gibt die Höhe des Gesamtschadens pro Jahr an, wobei dieser auch aus mehreren Einzelschäden bestehen kann.

Das Bilden von Beständen

Risiken werden zu Beständen zusammengefasst, wobei die einzelnen Risiken eines Bestandes eine gewisse Ähnlichkeit aufweisen sollten, dies kann beispielsweise dadurch gegeben sein, dass sie dem gleichen Versicherungszweig angehören (z. B. KFZ-Haftpflichtversicherung, Feuerversicherung) oder gleiche Merkmale aufweisen (z. B. Versicherungssumme bei der Feuerversicherung). Ideal wäre ein Bestand mit identischen Risiken, dies lässt sich aber vor allem in der Schadenversicherung praktisch kaum realisieren (vgl. [Sch02, S. 141]).

Mathematisch gesehen ist ein Bestand eine Teilmenge aller Risiken, die ein Versicherungsunternehmen übernommen hat. Mit dem Bilden von Beständen soll der Ausgleich im Kollektiv erreicht werden, das heißt, wenn in einem Bestand ein sehr hoher Schaden auftritt, dann wird er auf viele Risiken aufgeteilt. Große und kleine Schäden innerhalb eines Bestandes gleichen einander so weitgehend aus. Dass ein Bestand einerseits möglichst groß sein sollte, damit sich Schäden ausgleichen können, andererseits aber nur Risiken enthalten soll, die einander ähnlich sind, ist in gewisser Weise ein Widerspruch, für den praktisch ein Kompromiss gefunden werden muss (vgl. [Sch02, S. 141]).

Die mathematische Beschreibung der einzelnen Risiken und des Gesamtschadens

Jedes Risiko X_i wird durch eine Verteilungsfunktion $F_i(x)$ beschrieben:

$$F_i(x) = P(X_i \leq x).$$

Die Verteilungsfunktion gibt an, wie hoch die Wahrscheinlichkeit ist, dass die Zufallsvariable Werte annimmt, die nicht größer als x sind. (Für $x \leq 0$ gilt in der Versicherungsmathematik dabei stets $F(x) = 0$.) Falls nicht anders angegeben wird im Folgenden immer angenommen, dass alle Risiken eines Bestandes die gleiche Verteilung besitzen.

Da jedes Risiko die Höhe des Schadens pro Jahr angibt lässt sich die Zufallsvariable S für den jährlichen Gesamtschaden folgendermaßen berechnen:

$$S = \sum_{i=1}^n X_i.$$

Deshalb gilt für den Erwartungswert des jährlichen Gesamtschadens:

$$E(S) = \sum_{i=1}^n E(X_i).$$

Das heißt also, dass der zu erwartende Gesamtschaden pro Jahr einfach durch die Summe der zu erwartenden Einzelschäden berechnet werden kann. Dies entspricht dann auch dem Erwartungswert jenes Betrags, den das Versicherungsunternehmen an die VersicherungsnehmerInnen auszahlen muss. Allerdings ist der Erwartungswert alleine wenig aussagekräftig, für die Berechnung der Varianz kann die Formel

$$V(S) = \sum_{i=1}^n V(X_i)$$

nur verwendet werden, wenn die Risiken voneinander unabhängig sind. Da die Unabhängigkeit aber in manchen Fällen nicht vorausgesetzt werden kann sollte genau geprüft werden, ob die Risiken wirklich voneinander unabhängig sind. Risiken, bei denen die Unabhängigkeit nicht gegeben ist, treten beispielsweise auf, wenn sich mehrere Häuser aus einem Gebiet im Bestand befinden, die gegen Hagel- oder Hochwasserschäden versichert worden sind. Allgemein kann man sagen, dass die Unabhängigkeit immer dann nicht gegeben ist, wenn mehrere Risiken eines Bestandes der gleichen Gefahr ausgesetzt sind, in der Praxis kann es trotzdem vorkommen, dass auf die Unabhängigkeit verzichtet wird, dann liefert das Modell aber keine genaue Beschreibung des Bestandes, sondern kann ihn nur näherungsweise beschreiben (vgl. [Sch02, S. 144]).

Die Verteilungsfunktion F_S des jährlichen Gesamtschadens S ergibt sich bei unabhängigen Risiken aus der Faltung (vgl. [Sch95, S. 105]):

$$F_S(x) = P(S \leq x) = F_1(x) * F_2(x) * \dots * F_n(x).$$

Besitzen alle Risiken die gleiche Verteilung F ($F = F_1 = F_2 = \dots = F_n$), dann schreibt man:

$$F_S(x) = F^{*n}(x).$$

3.1.2 Daten zur Erstellung eines Modells

Eine Grundbedingung zur Durchführung versicherungstechnischer Rechnungen ist es, Informationen über den Schadenverlauf und die Art der einzelnen Risiken zu kennen.

Marktstatistiken, die die Daten mehrerer Versicherungsunternehmen umfassen, enthalten üblicherweise aber nur die Anzahl der Risiken, die Gesamtversicherungssumme und Daten über die Anzahl an Schäden und deren Summe, Daten über einzelne Risiken, wie deren Versicherungssumme und die Höhe einzelner Schäden, sind in diesen Statistiken nicht enthalten (vgl. [Mac02, S. 37]). Problematisch ist, dass sich einerseits die Größe eines Bestandes, andererseits seine Gesamtversicherungssumme über mehrere Jahre hinweg ändert, was zur Folge hat, dass sich auch die Verteilung des Gesamtschadens ändert. In der Praxis wird deswegen nicht der Gesamtschaden als maßgebliche Größe verwendet, sondern volumenabhängige Größen wie der Schadensatz oder der Schadenbedarf (vgl. [Mac02, S. 37]).

Der Schadenbedarf wird im Abschnitt über homogene Bestände behandelt, der Schadensatz im Abschnitt über inhomogene Bestände.

3.1.3 Homogene Bestände

Ein homogener Bestand kann folgendermaßen definiert werden:

Definition 3.2.

Ein Bestand von Risiken $\{X_1, X_2, \dots, X_n\}$ heißt homogen, wenn alle Risiken identisch verteilt und unabhängig sind.

Da in einem homogenen Bestand alle Risiken identisch verteilt sind und die gleiche Verteilung besitzen sind der Mittelwert und die Varianz der einzelnen Risiken konstant:

$$E(X_i) = m \quad \forall i = 1, 2, \dots, n;$$

$$V(X_i) = s^2 \quad \forall i = 1, 2, \dots, n.$$

Deshalb gilt für den Erwartungswert und die Varianz des Gesamtschadens:

$$E(S) = \sum_{i=1}^n E(X_i) = \sum_{i=1}^n m = n \cdot m,$$

$$V(S) = \sum_{i=1}^n V(X_i) = \sum_{i=1}^n s^2 = n \cdot s^2.$$

Schätzen der Parameter

Wie man erkennen kann hängt der Erwartungswert des Gesamtschadens in einem homogenen Bestand von der Größe des Bestandes ab. Da sich die Größe des Bestandes bei der Betrachtung mehrerer Jahre laufend ändert ist für Schätzungen in der Praxis eine Größe vorteilhaft, die nicht von der Größe des Bestandes abhängt. Diese Größe wird als Schadenbedarf bezeichnet und ist für einen homogenen Bestand mit n Risiken folgendermaßen definiert (vgl. [Mac02, S. 39]):

$$Z := \frac{S}{n} = \frac{1}{n} \sum_{i=1}^n X_i.$$

Für den Erwartungswert des Schadenbedarfs gilt aufgrund des Erwartungswertes für den Gesamtschaden:

$$E(Z) = m.$$

Deshalb ist der Schadenbedarf Z ein erwartungstreuer Schätzer für m . Weiters ergibt sich für die Varianz des Schadenbedarfs:

$$V(Z) = \frac{s^2}{n}.$$

An dieser Formel kann man einen Effekt des Ausgleichs im Kollektiv erkennen: Bildet man größere Bestände, dann nimmt die Varianz des Schadenbedarfs ab. Dadurch wird eine genauere Schätzung des Schadenbedarfs möglich.

Mit der Ungleichung von Tschebyschev erhält man folgendes Resultat für den Schadenbedarf in homogenen Beständen:

Satz 3.1. *Sei $\{X_1, X_2, \dots, X_n\}$ ein homogener Bestand von n Risiken. Dann gilt für den Schadenbedarf Z und für alle $c \in (0, \infty)$:*

$$P(|Z - m| \leq c) \geq 1 - \frac{s^2}{nc^2}.$$

Beweis. nach [Sch02, S. 150]

Für den Erwartungswert und die Varianz von Z gilt:

$$E(Z) = m,$$

$$V(Z) = \frac{s^2}{n}.$$

Die Ungleichung von Tschebyschev besagt, dass

$$P(|Z - E(Z)| \leq c) \geq 1 - \frac{V(Z)}{c^2} \quad \forall c \in (0, \infty).$$

Setzt man m für $E(Z)$ und $\frac{s^2}{n}$ für $V(Z)$ ein, dann ergibt sich:

$$P(|Z - m| \leq c) \geq 1 - \frac{s^2}{nc^2}.$$

□

Legt man eine obere Schranke $c \in \mathbb{R}, c > 0$ für den Schätzfehler fest, dann ergibt sich mit einer Fehlerwahrscheinlichkeit $\eta \in (0, 1), \eta \geq \frac{s^2}{nc^2}$, dass der Schätzfehler mit einer Wahrscheinlichkeit $1 - \eta$ kleiner als c oder gleich c ist (vgl. [Sch02, S.150]):

$$P(|Z - m| \leq c) \geq 1 - \eta.$$

Unter Verwendung des obigen Satzes erhält man außerdem folgendes Resultat:

$$\lim_{n \rightarrow \infty} P(|Z - m| \leq c) = 1.$$

Das heißt also, dass für den Grenzfall eines unendlich großen Bestandes die Differenz zwischen dem Schadenbedarf Z und dem erwarteten Schaden m eines Risikos fast sicher nicht größer als c ist (unabhängig vom Wert für c).

Zur Schätzung der Varianz kann der folgende Satz verwendet werden:

Satz 3.2. *Sei $\{X_1, X_2, \dots, X_n\}$ ein homogener Bestand von n Risiken mit $n \geq 2$ und sei Z der Schadenbedarf des Bestandes. Dann gilt:*

$$E \left(\frac{1}{n-1} \sum_{i=1}^n (X_i - Z)^2 \right) = s^2.$$

Beweis. nach [Sch02, S. 151]

Es gilt:

$$E \left(\frac{1}{n-1} \sum_{i=1}^n (X_i - Z)^2 \right) = \frac{1}{n-1} \cdot E \left(\sum_{i=1}^n (X_i - Z)^2 \right) =$$

$$= \frac{1}{n-1} \cdot \sum_{i=1}^n E((X_i - Z)^2).$$

Für alle $i \in \{1, 2, \dots, n\}$ gilt:

$$\begin{aligned} E((X_i - Z)^2) &= E(X_i^2 - 2X_iZ + Z^2) = \\ &= \underbrace{E(X_i^2)}_{=V(X_i)+E(X_i)^2} - 2E(X_iZ) + \underbrace{E(Z^2)}_{=V(Z)+E(Z)^2}. \end{aligned}$$

$E(X_iZ)$ kann folgendermaßen aufgeschrieben werden:

$$\begin{aligned} E(X_iZ) &= E\left(X_i \cdot \frac{1}{n} \sum_{j=1}^n X_j\right) = \\ &= \frac{1}{n} E(X_iX_1 + X_iX_2 + \dots + X_iX_i + \dots + X_iX_{n-1} + X_iX_n) \\ &= \frac{1}{n} E(X_iX_1) + E(X_iX_2) + \dots + E(X_iX_i) + \dots + E(X_iX_{n-1}) + E(X_iX_n). \end{aligned}$$

Da die Risiken $X_1, X_2, \dots, X_n$ unabhängig voneinander sind gilt:

$$E(X_jX_k) = E(X_j) \cdot E(X_k) \quad \forall j \neq k.$$

Das kann nun in den Ausdruck für $E(X_iZ)$ eingesetzt werden:

$$\begin{aligned} E(X_iZ) &= \frac{1}{n} \left(\underbrace{E(X_i)E(X_1)}_{=m \cdot m = m^2} + \underbrace{E(X_i)E(X_2)}_{=m^2} + \dots + \underbrace{E(X_iX_i)}_{=V(X_i)+E(X_i)^2} + \dots + \right. \\ &\quad \left. + \underbrace{E(X_iE(X_{n-1}))}_{=m^2} + \underbrace{E(X_iE(X_n))}_{=m^2} \right) = \\ &= \frac{1}{n} \left((n-1)m^2 + \underbrace{V(X_i)}_{=s^2} + \underbrace{E(X_i)^2}_{=m^2} \right) = \\ &= \frac{1}{n} (nm^2 + s^2) = \\ &= m^2 + \frac{s^2}{n}. \end{aligned}$$

Für $E((X_i - Z)^2)$ ergibt sich daraus:

$$\begin{aligned}
E((X_i - Z)^2) &= \underbrace{V(X_i)}_{=s^2} + \underbrace{E(X_i^2)}_{=m^2} - 2 \left(m^2 + \frac{s^2}{n} \right) + \underbrace{V(Z)}_{=\frac{s^2}{n}} + \underbrace{E(Z)^2}_{=m^2} = \\
&= s^2 - \frac{s^2}{n} = \\
&= \frac{n-1}{n} s^2.
\end{aligned}$$

Das kann schließlich in den ursprünglichen Ausdruck eingesetzt werden:

$$\begin{aligned}
E \left(\frac{1}{n-1} \sum_{i=1}^n (X_i - Z)^2 \right) &= \frac{1}{n-1} \cdot \underbrace{\sum_{i=1}^n \underbrace{E((X_i - Z)^2)}_{=\frac{n-1}{n} s^2}}_{=(n-1)s^2} = \\
&= s^2.
\end{aligned}$$

□

Da $E \left(\frac{1}{n-1} \sum_{i=1}^n (X_i - Z)^2 \right) = s^2$ gilt ist $\frac{1}{n-1} \sum_{i=1}^n (X_i - Z)^2$ ein erwartungstreuer Schätzer für die Varianz eines Risikos X_i .

Sind in einem Bestand $\{X_1, X_2, \dots, X_n\}$ von n Risiken deren Realisierungen $x_1, x_2, \dots, x_n$ am Ende einer Periode bekannt, dann können in der folgenden Periode die erwartungstreuen Schätzer $\hat{m}$ für m und $\hat{s}^2$ für s^2 verwendet werden (vgl. [Mac02, S. 39]):

$$\begin{aligned}
\hat{m} &= \frac{1}{n} \sum_{i=1}^n x_i, \\
\hat{s}^2 &= \frac{1}{n-1} \sum_{i=1}^n (x_i - \hat{m})^2.
\end{aligned}$$

Da wie bereits erläutert oftmals keine Daten für die einzelnen Risiken bekannt sind, sondern nur die Größe des Bestandes in mehreren Jahren und die jeweilige Realisierung des Schadenbedarfs, kann man auch diese Daten benutzen um erwartungstreue Schätzer für m und s^2 zu konstruieren: Seien $n_1, n_2, \dots, n_j$ die Größen des Bestandes in j betrachteten Perioden und sei z_i die (inflationsbereinigte) Realisierung des Schadenbedarfs in der i . Periode ($i \in \{1, 2, \dots, j\}$). Dann sind $\hat{m}$ und $\hat{s}^2$ erwartungstreue Schätzer für m

bzw. s^2 (vgl. [Mac02, S. 39]):

$$\hat{m} = \frac{1}{\sum_{i=1}^j n_i} \sum_{i=1}^j n_i z_i,$$
$$\hat{s}^2 = \frac{1}{j-1} \sum_{i=1}^j n_i (z_i - \hat{m})^2.$$

3.1.4 Die Verteilung der Schadenshöhen und die Verteilung des Gesamtschadens

Bisher wurden die einzelnen Risiken und der Gesamtschaden vor allem durch ihren Erwartungswert und ihre Varianz charakterisiert. In diesem Abschnitt wird darauf eingegangen, welche Verteilungen zur Beschreibung von Risiken bzw. des Gesamtschadens geeignet sind und wie man aus der Verteilung einzelner Risiken die Verteilung des Gesamtschadens ermitteln kann.

Nicht jede beliebige Verteilung eignet sich zur Beschreibung von Risiken. Um ein passendes Modell für die Verteilung der Schadenshöhen zu finden muss die Verteilung eine gewisse Gestalt besitzen, die folgende Eigenschaften aufweist (nach [5, S. 5]):

- Schäden mit einer geringen Schadenshöhe (sogenannte „Bagatellschäden“) treten nur selten auf. (Ein Grund dafür ist, dass kleine Schäden dem Versicherungsunternehmen nicht gemeldet werden, weil die Reperatur- oder Regulierungskosten nur gering sind bzw. weil VersicherungsnehmerInnen negative Konsequenzen, beispielsweise die Einstufung in einer ungünstigeren Stufe in einem Bonus-/Malus-System drohen.)
- Mittlere Schadenshöhen kommen häufig vor.
- Sehr hohe Schäden treten wie geringe Schäden nur selten auf.

Außerdem muss bei der Modellierung beachtet werden, dass der weitaus größte Teil der Risiken keinen Schaden verursacht.

Es ist aufgrund der genannten Eigenschaften nicht zielführend, eine beliebige Verteilung anzusetzen und die Parameter so zu wählen, dass der Erwartungswert und die Varianz der Verteilung dem Erwartungswert und der Varianz der Risiken entsprechen, denn der Erwartungswert und die Varianz sagen wenig über die Gestalt der Verteilung aus (vgl. [Mac02, S. 46]). (Viele Verteilungen könnten so angepasst werden, dass ihr

Erwartungswert und die Varianz übereinstimmen, die Gestalt würde sich aber stark unterscheiden.)

Obwohl die Höhe eines Schadens nur bestimmte diskrete Werte (nämlich Geldbeträge) annehmen kann ist es üblich, Schadenshöhen mit kontinuierlichen Verteilungen zu beschreiben. Da hohe Schäden nur selten auftreten sind für die Modellierung von Schadenshöhen außerdem nur rechtsschiefe Verteilungen geeignet. Aufgrund der Vielzahl an Faltungen, die bei der Berechnung der Gesamtschadenverteilung auftreten, ist es aus rechentechnischen Gründen zudem sinnvoll eine Verteilung zu wählen, die eine einfache Berechnung der Faltung erlaubt.

Verteilungen, die sich für die Modellierung der Schadenshöhen einzelner Risiken eignen und die in den folgenden Abschnitten beschrieben werden sind:

- Gammaverteilung,
- Exponentialverteilung,
- Inverse Normalverteilung,
- Paretoverteilung.

Außerdem wird beschrieben, wie der Gesamtschaden eines Bestandes modelliert werden kann, ohne die Verteilung über die Faltung der Einzelschadenverteilungen zu berechnen. Als Verteilung für den Gesamtschaden kommt dabei die Logarithmische Normalverteilung zum Einsatz.

Gammaverteilung

Die Gammaverteilung ist eine zweiparametrische Verteilung, die den großen Vorteil besitzt, dass eine endliche Summe unabhängiger, gammaverteilter Zufallsvariablen wiederum gammaverteilt ist. Dadurch lässt sich die Verteilung des Gesamtschadens in einem Bestand von gammaverteilten Risiken explizit berechnen. Ein Nachteil für die numerische Berechnung besteht allerdings darin, dass zur Angabe der Dichte und der Verteilungsfunktion die Gammafunktion notwendig ist, weshalb diese nicht integralfrei angegeben werden können.

Da die Gammaverteilung eine zentrale Rolle in der Versicherungsmathematik spielt wird sie in diesem Abschnitt ausführlich behandelt und es werden wichtige Eigenschaften der Verteilung besprochen und hergeleitet.

Zur Beschreibung der Gammaverteilung wird die Gammafunktion benötigt. Diese ist folgendermaßen definiert:

Definition 3.3.

Die Abbildung

$$\Gamma : (0, \infty) \rightarrow \mathbb{R}, \Gamma(t) := \int_0^\infty x^{t-1} e^{-x} dx$$

wird als Gammafunktion bezeichnet.

Bevor die Gammaverteilung definiert wird werden noch zwei wichtige Eigenschaften der Gammafunktion gezeigt, die im Folgenden benötigt werden. Die erste Eigenschaft ist die Möglichkeit der rekursiven Berechnung der Gammafunktion:

Lemma 3.1. *Für die Gammafunktion gilt:*

$$\Gamma(t+1) = t\Gamma(t) \quad \forall t \in \mathbb{R}, t > 0.$$

Beweis.

Es gilt:

$$\begin{aligned} \Gamma(t+1) &= \int_0^\infty x^t e^{-x} dx = \\ &\quad (\text{partielle Integration: } u(x) = x^t, v'(x) = e^{-x}) \\ &= \underbrace{-x^t e^{-x}}_{=0} \Big|_0^\infty + \underbrace{\int_0^\infty t x^{t-1} e^{-x} dx}_{=t\Gamma(t)} = \\ &= t\Gamma(t). \end{aligned}$$

□

Die zweite Eigenschaft betrifft die Betrachtung der Gammafunktion auf den natürlichen Zahlen:

Lemma 3.2. *Sei $t \in \mathbb{N}$, dann gilt:*

$$\Gamma(t) = (t-1)! \quad \forall t \in \mathbb{N}, t \geq 1.$$

Beweis.

Dies lässt sich mit einer vollständigen Induktion beweisen:

Induktionsanfang: Für $t = 1$ gilt:

$$\Gamma(1) := \int_0^\infty e^{-x} dx = -e^{-x} \Big|_0^\infty = 1.$$

Induktionsvoraussetzung: Die Behauptung gelte für ein $n \in \mathbb{N}, n > 1$.

Induktionsschritt: Für $n + 1$ gilt nach dem obigen Lemma und nach der Induktionsvoraussetzung:

$$\Gamma(n + 1) = n \cdot \Gamma(n) = n!.$$

□

Die Gammafunktion kann also als Interpolation der Fakultätsfunktion auf den positiven reellen Zahlen betrachtet werden.

Nun kann die Gammaverteilung definiert werden:

Definition 3.4.

Als Gammaverteilung $\Gamma(\alpha, \beta)$ wird die Wahrscheinlichkeitsverteilung mit der Dichte

$$f(x) = \begin{cases} \frac{\beta^\alpha}{\Gamma(\alpha)} x^{\alpha-1} e^{-\beta x} & x > 0, \\ 0 & x \leq 0, \end{cases}$$

bezeichnet. Die Parameter $\alpha \in \mathbb{R}$ und $\beta \in \mathbb{R}$ müssen $\alpha > 0$ und $\beta > 0$ erfüllen.

Der Parameter β wird als Skalenparameter bezeichnet, er gibt an, wie stark die Verteilungsfunktion parallel zur Abszisse gestreckt (für $\beta < 1$) oder gestaucht (für $\beta > 1$) wird. Die Form der Verteilungsfunktion bzw. deren Dichte wird durch den Formenparameter α festgelegt (vgl. [5, S. 5]).

Für versicherungsmathematische Anwendungen sind vor allem Verteilungen mit $\alpha < 1$ geeignet. Der Grund dafür ist, dass der größte Teil der Risiken eines Bestandes keine Schäden verursacht.

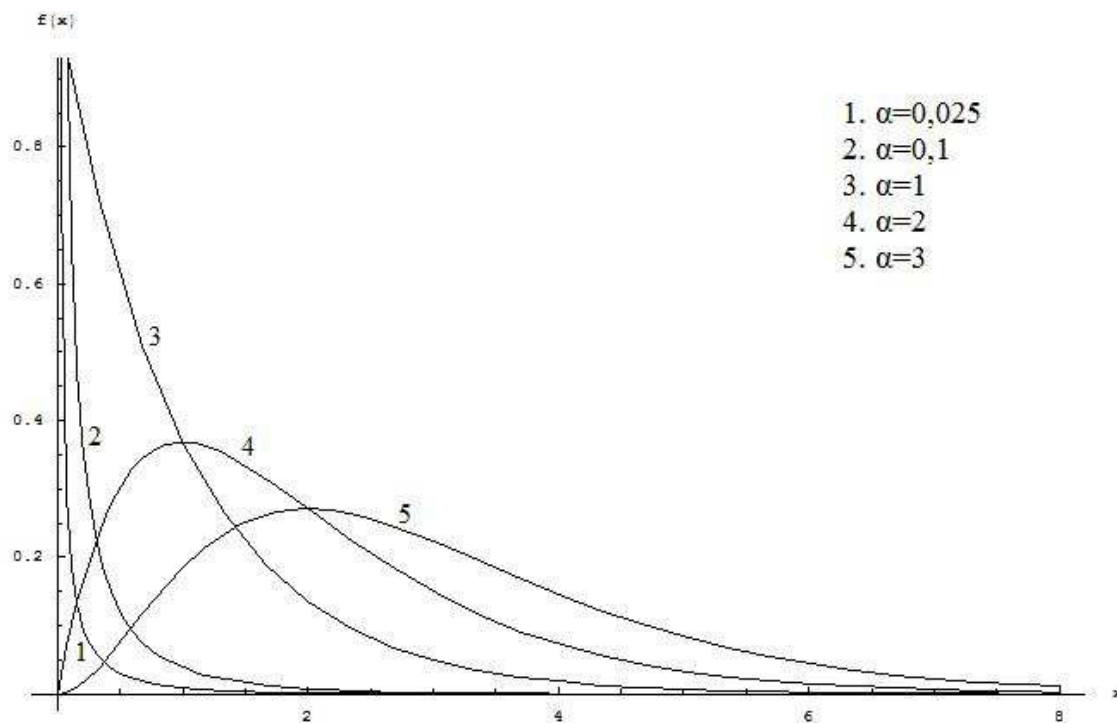


Abbildung 3.1: Die Dichtefunktion der Gammaverteilung mit $\beta = 1$ und verschiedenen Werten für α .

Bemerkung. Die Gammaverteilung steht in einem engen Zusammenhang mit der Exponential-, der Erlang- und der χ^2 -Verteilung:

- Ist $\alpha = 1$, dann entspricht die Gammaverteilung einer Exponentialverteilung mit dem Parameter β , es gilt:

$$\Gamma(1, \beta) = \text{Exp}(\beta) \quad \beta > 0.$$

- Ist $\alpha \in \mathbb{N}$, dann entspricht die Gammaverteilung einer Erlangverteilung mit den Parametern β und α , es gilt:

$$\Gamma(\alpha, \beta) = \text{Erl}(\beta, \alpha) \quad \alpha \in \mathbb{N}, \beta > 0.$$

- Ist $\alpha = \frac{n}{2}$ für ein $n \in \mathbb{N}$ und $\beta = \frac{1}{2}$, dann entspricht die Gammaverteilung einer χ^2 -Verteilung mit n Freiheitsgraden.

Um die Verteilungsfunktion der Gammaverteilung anzugeben wird die unvollständige Gammafunktion der oberen Grenze benötigt, die folgendermaßen definiert ist:

Definition 3.5.

Die Abbildung

$$\Gamma : (0, \infty) \times (0, \infty) \rightarrow \mathbb{R}, \Gamma(b, t) := \int_0^b x^{t-1} e^{-x} dx$$

wird als unvollständige Gammafunktion der oberen Grenze bezeichnet.

Satz 3.3. *Sei X eine gammaverteilte Zufallsvariable mit den Parametern α und β . Dann lautet die Verteilungsfunktion von X :*

$$F(x) = \begin{cases} \frac{\Gamma(\beta x, \alpha)}{\Gamma(\alpha)} & x > 0, \\ 0 & x \leq 0. \end{cases}$$

Beweis. nach [Sch95, S. 33]

Sei $x > 0$, dann gilt:

$$\begin{aligned} F(x) &= \int_{-\infty}^x f(t) dt = \\ &= \int_0^x \frac{\beta^\alpha}{\Gamma(\alpha)} t^{\alpha-1} e^{-\beta t} dt = \\ &= \frac{\beta^\alpha}{\Gamma(\alpha)} \int_0^x t^{\alpha-1} e^{-\beta t} dt = \\ &\text{(Substitution: } u = \beta t) \\ &= \frac{\beta^\alpha}{\Gamma(\alpha)} \int_0^{\beta x} \left(\frac{u}{\beta}\right)^{\alpha-1} e^{-u} \frac{1}{\beta} du = \\ &= \frac{1}{\Gamma(\alpha)} \underbrace{\int_0^{\beta x} u^{\alpha-1} e^{-u} du}_{=\Gamma(\beta x, \alpha)} = \\ &= \frac{\Gamma(\beta x, \alpha)}{\Gamma(\alpha)}. \end{aligned}$$

Sei $x \leq 0$, dann gilt:

$$F(x) = 0.$$

Insgesamt gilt daher:

$$F(x) = \begin{cases} \frac{\Gamma(\beta x, \alpha)}{\Gamma(\alpha)} & x > 0, \\ 0 & x \leq 0. \end{cases}$$

□

Satz 3.4. Für eine gammaverteilte Zufallsvariable X mit den Parametern α und β gilt:

$$E(X) = \frac{\alpha}{\beta},$$

$$V(X) = \frac{\alpha}{\beta^2}.$$

Beweis.

Für $E(X)$ gilt:

$$\begin{aligned} E(X) &= \int_{-\infty}^{\infty} x f(x) dx = \\ &= \int_0^{\infty} x \frac{\beta^\alpha}{\Gamma(\alpha)} x^{\alpha-1} e^{-\beta x} dx = \\ &= \frac{\beta^\alpha}{\Gamma(\alpha)} \int_0^{\infty} x^\alpha e^{-\beta x} dx = \\ &\quad (\text{Substitution: } u = \beta x) \\ &= \frac{\beta^\alpha}{\Gamma(\alpha)} \int_0^{\infty} \left(\frac{u}{\beta}\right)^\alpha e^{-u} \frac{1}{\beta} du = \\ &= \frac{1}{\Gamma(\alpha)\beta} \underbrace{\int_0^{\infty} u^\alpha e^{-u} du}_{=\Gamma(\alpha+1)=\alpha\Gamma(\alpha)} = \\ &= \frac{\alpha}{\beta}. \end{aligned}$$

Für $E(X^2)$ gilt:

$$\begin{aligned} E(X^2) &= \int_{-\infty}^{\infty} x^2 f(x) dx = \\ &= \int_0^{\infty} x^2 \frac{\beta^\alpha}{\Gamma(\alpha)} x^{\alpha-1} e^{-\beta x} dx = \\ &= \frac{\beta^\alpha}{\Gamma(\alpha)} \int_0^{\infty} x^{\alpha+1} e^{-\beta x} dx = \\ &\quad (\text{Substitution: } u = \beta x) \end{aligned}$$

$$\begin{aligned}
&= \frac{\beta^\alpha}{\Gamma(\alpha)} \int_0^\infty \left(\frac{u}{\beta}\right)^{\alpha+1} e^{-u} \frac{1}{\beta} du = \\
&= \frac{1}{\Gamma(\alpha)\beta^2} \underbrace{\int_0^\infty u^{\alpha+1} e^{-u} du}_{=\Gamma(\alpha+2)=(\alpha+1)\Gamma(\alpha+1)=(\alpha+1)\alpha\Gamma(\alpha)} \\
&= \frac{\alpha(\alpha+1)}{\beta^2}.
\end{aligned}$$

Daraus ergibt sich für $V(X)$:

$$\begin{aligned}
V(X) &= E(X^2) - (E(X))^2 = \\
&= \frac{\alpha(\alpha+1)}{\beta^2} - \left(\frac{\alpha}{\beta}\right)^2 = \\
&= \frac{\alpha^2 + \alpha - \alpha^2}{\beta^2} = \\
&= \frac{\alpha}{\beta^2}.
\end{aligned}$$

□

Da in der Versicherungsmathematik der Erwartungswert eine große Rolle spielt, wird eine Darstellung der Dichte verwendet, in der der Erwartungswert direkt vorkommt: Setzt man $\mu := E(X) = \frac{\alpha}{\beta}$, dann hat die Dichtefunktion folgende Darstellung (vgl. [10, S. 22]):

$$f(x) = \begin{cases} \frac{(\frac{\alpha}{\mu})^\alpha}{\Gamma(\alpha)} x^{\alpha-1} e^{-\frac{\alpha}{\mu}x} & x > 0, \\ 0 & x \leq 0. \end{cases}$$

Um den nachfolgenden Satz (Berechnung der Faltung) zu beweisen wird folgendes Resultat benötigt:

Lemma 3.3. *Es gilt:*

$$\int_0^1 (1-x)^{\alpha-1} x^{\beta-1} dx = \frac{\Gamma(\alpha)\Gamma(\beta)}{\Gamma(\alpha+\beta)} \quad \forall \alpha > 0, \beta > 0.$$

Eine Begründung dieses Resultats findet man beispielsweise hier: [Haf89, S. 72].

Für zwei unabhängige, gammaverteilte Zufallsvariablen lässt sich die Faltung explizit berechnen:

Satz 3.5. *Seien X_1 und X_2 zwei unabhängige Zufallsvariablen, für die gilt: X_1 ist $\Gamma(\alpha_1, \beta)$ -verteilt und X_2 ist $\Gamma(\alpha_2, \beta)$ -verteilt. Dann gilt: Die Zufallsvariable $X_1 + X_2$ ist $\Gamma(\alpha_1 + \alpha_2, \beta)$ -verteilt.*

Beweis. nach [5, S. 23]

Sei f_1 die Dichte der Zufallsvariable X_1 und f_2 die Dichte der Zufallsvariable X_2 , dann gilt für die Dichte g von $X_1 + X_2$:

$$g(x) = \int_{-\infty}^{\infty} f_1(x-t)f_2(t)dt.$$

Für $x > 0$ gilt:

$$\begin{aligned} g(x) &= \int_0^x \frac{\beta^{\alpha_1}}{\Gamma(\alpha_1)}(x-t)^{\alpha_1-1}e^{-\beta(x-t)} \frac{\beta^{\alpha_2}}{\Gamma(\alpha_2)}t^{\alpha_2-1}e^{-\beta t}dt = \\ &= \frac{\beta^{\alpha_1+\alpha_2}}{\Gamma(\alpha_1)\Gamma(\alpha_2)} \int_0^x (x-t)^{\alpha_1-1}e^{-\beta x}t^{\alpha_2-1}dt = \\ &\text{(Substitution: } u = \frac{t}{x} \Rightarrow t = ux; \frac{dt}{du} = x) \\ &= \frac{\beta^{\alpha_1+\alpha_2}}{\Gamma(\alpha_1)\Gamma(\alpha_2)}e^{-\beta x} \int_0^1 \underbrace{(x-ux)^{\alpha_1-1}}_{=x^{\alpha_1-1}(1-u)^{\alpha_1-1}} (ux)^{\alpha_2-1}xdu = \\ &= \frac{\beta^{\alpha_1+\alpha_2}}{\Gamma(\alpha_1)\Gamma(\alpha_2)}e^{-\beta x}x^{\alpha_1+\alpha_2-1} \underbrace{\int_0^1 (1-u)^{\alpha_1-1}u^{\alpha_2-1}du}_{=\frac{\Gamma(\alpha_1)\Gamma(\alpha_2)}{\Gamma(\alpha_1+\alpha_2)}} = \\ &= \frac{\beta^{\alpha_1+\alpha_2}}{\Gamma(\alpha_1+\alpha_2)}e^{-\beta x}x^{\alpha_1+\alpha_2-1}. \end{aligned}$$

Für $x \leq 0$ gilt:

$$g(x) = 0.$$

Insgesamt gilt daher:

$$g(x) = \begin{cases} \frac{\beta^{\alpha_1+\alpha_2}}{\Gamma(\alpha_1+\alpha_2)}x^{\alpha_1+\alpha_2-1}e^{-\beta x} & x > 0, \\ 0 & x \leq 0. \end{cases}$$

Die Zufallsvariable $X_1 + X_2$ ist also $\Gamma(\alpha_1 + \alpha_2, \beta)$ -verteilt. $\square$

Der obige Satz lässt sich auf eine endliche Folge von unabhängigen Zufallsvariablen erweitern:

Korollar. Sei $X_1, X_2, \dots, X_n$ eine Folge von unabhängigen Zufallsvariablen, für die gilt: X_i ist $\Gamma(\alpha_i, \beta)$ -verteilt, $i = 1, 2, \dots, n$. Die Summe $\sum_{k=1}^n X_k$ ist dann $\Gamma(\sum_{k=1}^n \alpha_k, \beta)$ -verteilt.

Beweis.

Dies lässt sich mit einer vollständigen Induktion beweisen:

Für $n = 2$ gilt die Behauptung nach dem obigen Satz.

Induktionsvoraussetzung: Die Behauptung gelte für eine Folge von $n - 1$ unabhängigen, gammaverteilten Zufallsvariablen.

Induktionsschritt: Da nach der Induktionsvoraussetzung für die Zufallsvariable $\sum_{k=1}^{n-1} X_k$ gilt, dass sie $\Gamma(\sum_{k=1}^{n-1} \alpha_k, \beta)$ -verteilt ist und die Zufallsvariable X_n nach Voraussetzung $\Gamma(\alpha_n, \beta)$ -verteilt ist und da $\sum_{k=1}^{n-1} X_k$ und X_n nach Voraussetzung unabhängig sind, lässt sich wiederum der obige Satz anwenden und es gilt: Die Zufallsvariable $\sum_{k=1}^n X_k$ ist $\Gamma(\sum_{k=1}^n \alpha_k, \beta)$ -verteilt. $\square$

Für einen Bestand von gammaverteilten Risiken heißt das, dass auch der Gesamtschaden gammaverteilt ist. Betrachtet man einen homogenen Bestand von n Risiken, die alle jeweils $\Gamma(\alpha, \beta)$ -verteilt sind, dann gilt für den Gesamtschaden S : S ist $\Gamma(n\alpha, \beta)$ -verteilt. Deshalb gilt für die Dichte und für die Verteilungsfunktion des Gesamtschadens:

$$f_S(x) = \begin{cases} \frac{\beta^{n\alpha}}{\Gamma(n\alpha)} x^{n\alpha-1} e^{-\beta x} & x > 0, \\ 0 & x \leq 0, \end{cases}$$

$$F_S(x) = \begin{cases} \frac{\Gamma(\beta x, n\alpha)}{\Gamma(n\alpha)} & x > 0, \\ 0 & x \leq 0. \end{cases}$$

Außerdem gilt für den Erwartungswert und die Varianz von S :

$$E(S) = n \frac{\alpha}{\beta},$$

$$V(S) = n \frac{\alpha}{\beta^2}.$$

Exponentialverteilung

Im Gegensatz zur Gammaverteilung ist die Exponentialverteilung eine einparametrische Verteilung. Das wirkt sich nachteilig aus, denn verglichen mit der Gammaverteilung ist sie dadurch nicht so gut anpassbar und kann mitunter nur grobe Näherungen liefern. Demgegenüber steht aber die Einfachheit von Berechnungen, denn im Gegensatz zur Gammaverteilung lassen sich die Dichte und die Verteilungsfunktion integralfrei angeben. Aufgrund der rechentechnischen Vorteile wird die Exponentialverteilung deshalb trotzdem als Verteilung von Schadenshöhen verwendet.

Wie schon beschrieben kann die Exponentialverteilung als Gammaverteilung mit dem Formenparameter $\alpha = 1$ betrachtet werden, die Resultate der Gammaverteilung können deshalb zur Beschreibung der Exponentialverteilung benutzt werden.

Definition 3.6.

Als Exponentialverteilung $\text{Exp}(\lambda)$ wird die Wahrscheinlichkeitsverteilung mit der Dichte

$$f(x) = \begin{cases} \lambda e^{-\lambda x} & x > 0, \\ 0 & x \leq 0, \end{cases}$$

bezeichnet. Der Parameter λ muss $\lambda > 0$ erfüllen.

Für die Verteilungsfunktion, den Erwartungswert und die Varianz gelten aufgrund der Eigenschaften der Gammaverteilung:

$$F(x) = \begin{cases} 1 - e^{-\lambda x} & x > 0, \\ 0 & x \leq 0, \end{cases}$$

$$E(X) = \frac{1}{\lambda},$$

$$V(X) = \frac{1}{\lambda^2}.$$

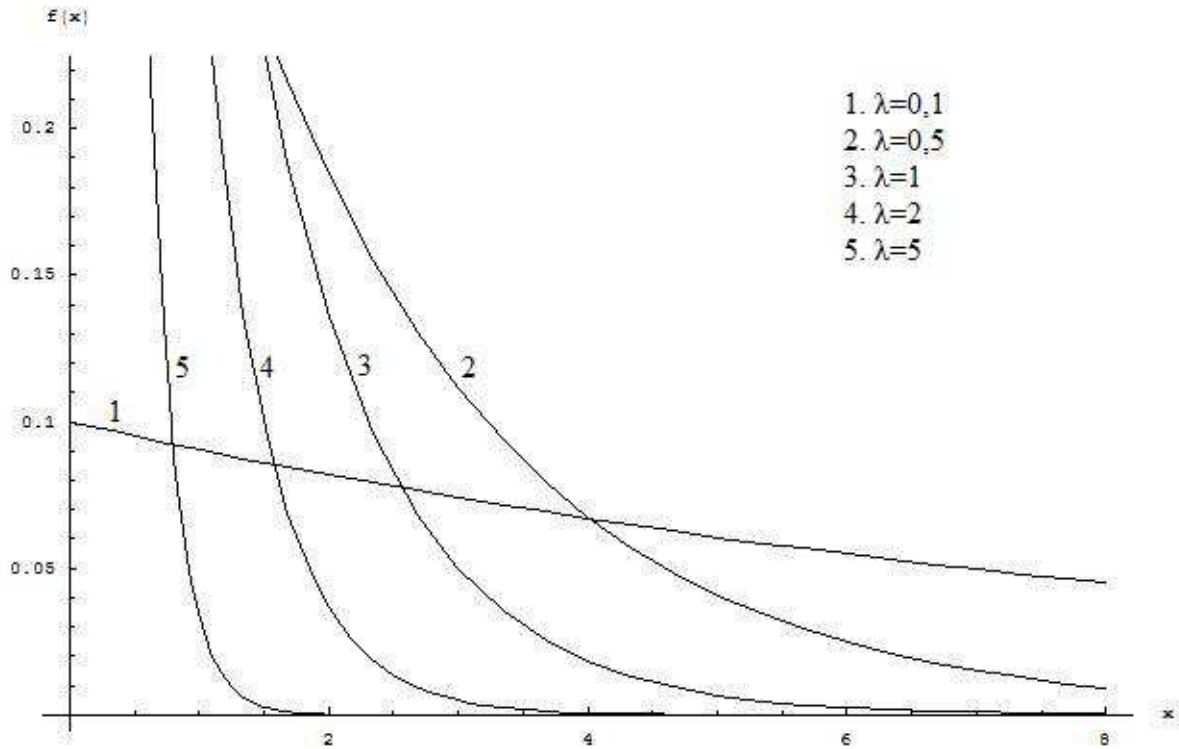


Abbildung 3.2: Die Dichtefunktion der Exponentialverteilung mit verschiedenen Werten für λ .

Betrachtet man eine Folge $X_1, X_2, \dots, X_n$ von $\text{Exp}(\lambda)$ -verteilten Zufallsvariablen, dann ist die Summe $\sum_{i=1}^n X_n$ gammaverteilt mit dem Formparameter n und dem Skalenparameter λ , d.h. $\Gamma(n, \lambda)$ -verteilt. Da $\Gamma(n) = (n-1)! \forall n \in \mathbb{N}$ lässt sich die Dichte von $\sum_{i=1}^n X_n$ im Gegensatz zur Gammaverteilung integralfrei angeben.

Für einen homogenen Bestand von $\text{Exp}(\lambda)$ -verteilten Risiken $X_1, X_2, \dots, X_n$ ergibt sich deshalb für die Dichte und die Verteilungsfunktion des Gesamtschadens S :

$$f_S(x) = \begin{cases} \frac{\lambda^n}{(n-1)!} x^{n-1} e^{-\lambda x} & x > 0, \\ 0 & x \leq 0, \end{cases}$$

$$F_S(x) = \begin{cases} \frac{\Gamma(\lambda x, n)}{(n-1)!} & x > 0, \\ 0 & x \leq 0. \end{cases}$$

Weiters gilt für den Erwartungswert und die Varianz des Gesamtschadens:

$$E(S) = \frac{n}{\lambda},$$

$$V(S) = \frac{n}{\lambda^2}.$$

Inverse Normalverteilung

Nach der ausführlichen Behandlung der Gamma- und der Exponentialverteilung sollen nun noch die Eigenschaften von zwei weiteren wichtigen Verteilungen für die Versicherungsmathematik, der Inversen Normalverteilung und der Paretoverteilung, angeführt werden.

Die Inverse Normalverteilung (auch als Inverse Gaußverteilung bezeichnet) ist wie die Gammaverteilung eine zweiparametrische Verteilung. Analog zur Gammaverteilung kann die Faltung auch bei der Inversen Normalverteilung explizit berechnet werden, außerdem bietet die Inverse Normalverteilung den Vorteil, dass die Dichte und die Verteilungsfunktion integralfrei angegeben werden können. (Die für die Inverse Normalverteilung angegebenen Daten wurden entnommen aus [Sch95, S. 46ff].) Die Definition der Inversen Normalverteilung lautet folgendermaßen:

Definition 3.7.

Als Inverse Normalverteilung $IN(\lambda, \mu)$ wird die Wahrscheinlichkeitsverteilung mit der Dichte

$$f(x) = \begin{cases} \sqrt{\frac{\lambda}{2\pi x^3}} \cdot e^{-\frac{\lambda(x-\mu)^2}{2\mu^2 x}} & x > 0, \\ 0 & x \leq 0, \end{cases}$$

bezeichnet. Die Parameter $\lambda \in \mathbb{R}$ und $\mu \in \mathbb{R}$ müssen $\lambda > 0$ und $\mu > 0$ erfüllen.

Der Parameter λ ist dabei ein Formparameter und wird als Ereignisrate, der Parameter μ als Mittelwert bezeichnet. (Diese Bezeichnungen kommen daher, dass die Inverse Normalverteilung ursprünglich zur Untersuchung der Brownschen Bewegung entwickelt und verwendet wurde (vgl. [1]).)

Sei X eine $IN(\lambda, \mu)$ -verteilte Zufallsvariable, dann gilt:

$$E(X) = \mu,$$

$$V(X) = \frac{\mu^3}{\lambda}.$$

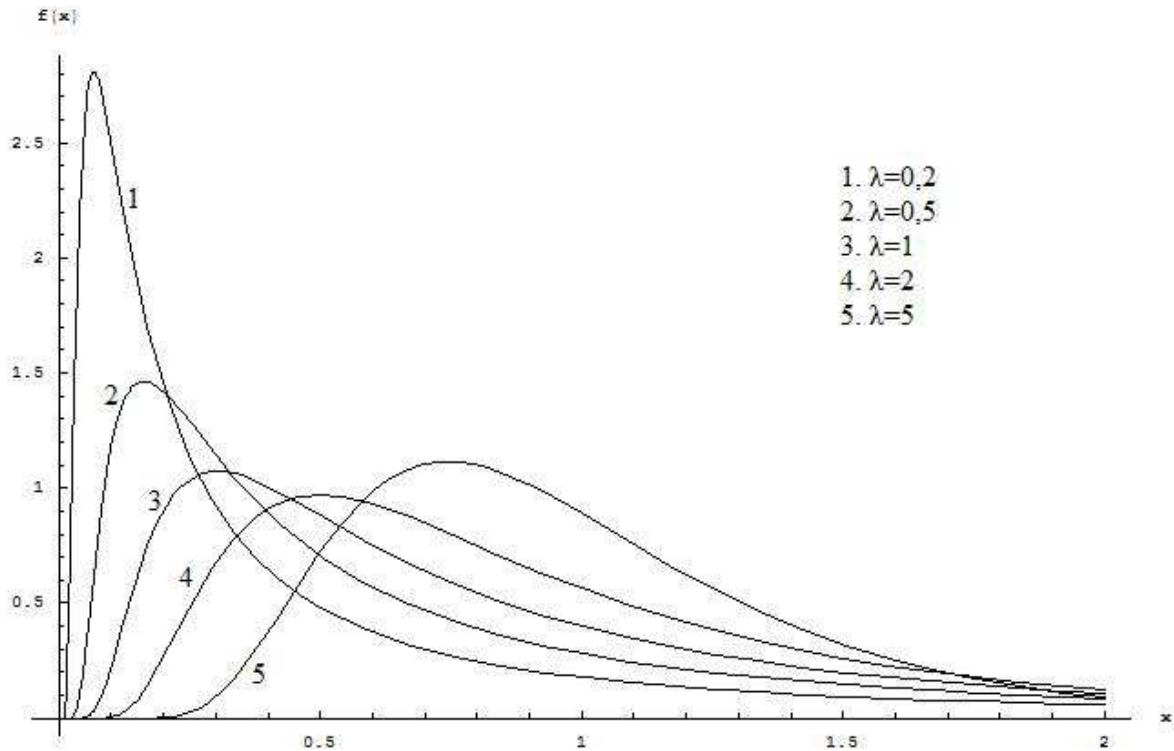


Abbildung 3.3: Die Dichtefunktion der Inversen Normalverteilung mit $\mu = 1$ und verschiedenen Werten für λ .

Die Verteilungsfunktion der inversen Normalverteilung kann folgendermaßen angegeben werden:

$$F(x) = \begin{cases} \Phi\left(\sqrt{\frac{\lambda}{x}} \cdot \left(\frac{x}{\mu} - 1\right)\right) + e^{\frac{2\lambda}{\mu}} \Phi\left(-\sqrt{\frac{\lambda}{x}} \cdot \left(\frac{x}{\mu} + 1\right)\right) & x > 0, \\ 0 & x \leq 0. \end{cases}$$

Dabei bezeichnet $\Phi(x)$ die Verteilungsfunktion der Standardnormalverteilung (0-1-Normalverteilung).

Für Anwendungen in der Versicherungsmathematik ist folgendes Resultat von besonderer Bedeutung: Sei $X_1, X_2, \dots, X_n$ eine Folge von unabhängigen, invers-normalverteilten Zufallsvariablen mit den Parametern λ und μ . Dann ist die Zufallsvariable $\sum_{i=1}^n X_i$ ebenfalls invers-normalverteilt mit den Parametern $n^2\lambda$ und $n\mu$.

Umgelegt auf einen homogenen Bestand von n invers-normalverteilten Risiken mit den Parametern λ und μ heißt das, dass der Gesamtschaden ebenfalls invers-normalverteilt mit den Parametern $n^2\lambda$ und $n\mu$ ist und dass für den Erwartungswert bzw. die Varianz

des Gesamtschadens gilt:

$$E(S) = n\mu,$$

$$V(S) = \frac{n^3\mu^3}{n^2\lambda} = \frac{n\mu^3}{\lambda}.$$

Außerdem gelten für die Dichte und für die Verteilungsfunktion des Gesamtschadens:

$$f_S(x) = \begin{cases} \sqrt{\left(\frac{n^2\lambda}{2\pi x^3}\right)} \cdot e^{-\frac{\lambda(x-n\mu)^2}{2\mu^2 x}} & x > 0, \\ 0 & x \leq 0, \end{cases}$$

$$F_S(x) = \begin{cases} \Phi\left(\sqrt{\frac{n^2\lambda}{x}} \cdot \left(\frac{x}{n\mu} - 1\right)\right) + e^{\frac{2n\lambda}{\mu}} \Phi\left(-\sqrt{\frac{n^2\lambda}{x}} \cdot \left(\frac{x}{n\mu} + 1\right)\right) & x > 0, \\ 0 & x \leq 0. \end{cases}$$

Paretoverteilung

Die Paretoverteilung ist wie die Gammaverteilung und die Inverse Normalverteilung eine zweiparametrische Verteilung. Obwohl sie bezüglich der Faltung nicht so schöne Eigenschaften besitzt wie die bisher besprochenen Verteilungen wird sie dennoch in vielen Versicherungssparten eingesetzt (wie beispielsweise der Feuerversicherung, vgl. [Sch95, S. 35]). Der Grund dafür ist, dass die Dichte der Paretoverteilung im Gegensatz zu den bisher behandelten Verteilungen für große Werte sehr langsam gegen 0 konvergiert. Das heißt für das Erstellen von versicherungsmathematischen Modellen, dass hohe Schäden bei paretoverteilten Risiken wahrscheinlicher sind als bei den bisher behandelten Verteilungen. Deswegen eignet sich die Paretoverteilung vor allem für Risiken, bei denen hohe Schäden auftreten können. Für Risiken, bei denen sehr hohe Schäden auszuschließen sind ist die Paretoverteilung hingegen nicht geeignet (vgl. [Sch95, S. 35]). (Die Eignung der Paretoverteilung für Großschäden wird im nächsten Abschnitt noch genauer besprochen.)

Die Paretoverteilung kann folgendermaßen definiert werden:

Definition 3.8.

Als Paretoverteilung $\text{Par}(a, b)$ wird die Wahrscheinlichkeitsverteilung mit der Dichte

$$f(x) = \begin{cases} \frac{b}{a} \left(\frac{a}{a+x}\right)^{b+1} & x > 0, \\ 0 & x \leq 0, \end{cases}$$

bezeichnet. Die Parameter $a \in \mathbb{R}$ und $b \in \mathbb{R}$ müssen $a > 0$ und $b > 0$ erfüllen.

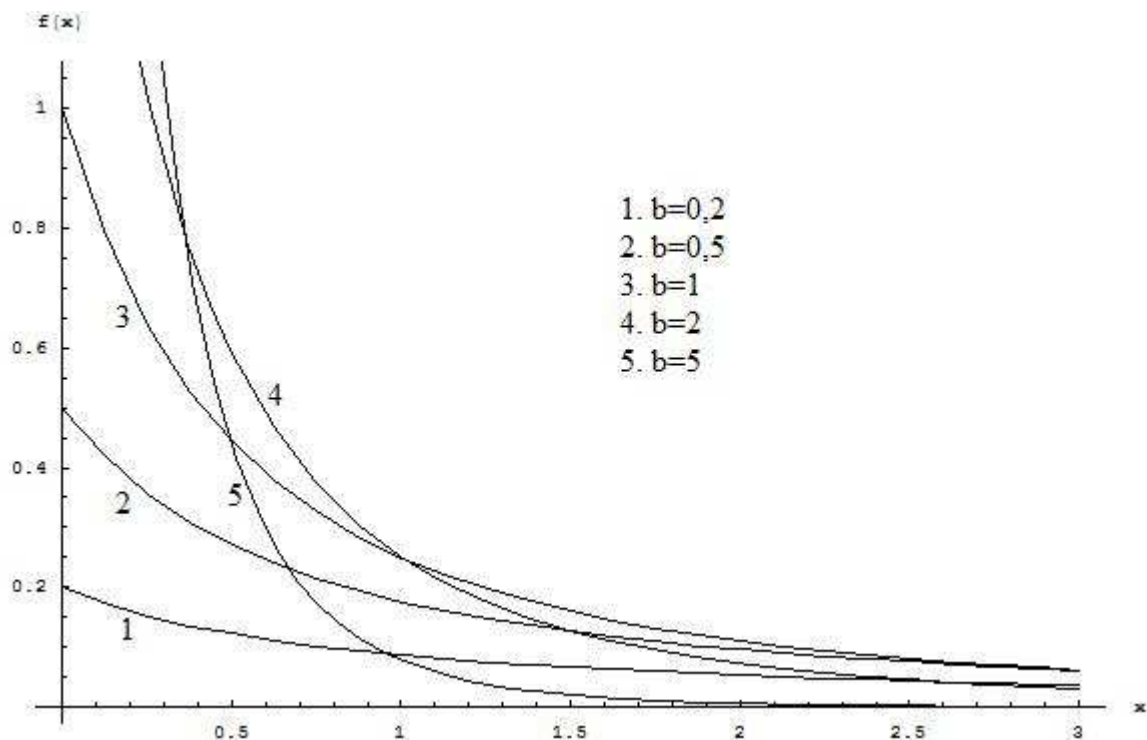


Abbildung 3.4: Die Dichtefunktion der Paretoverteilung mit $a = 1$ und verschiedenen Werten für b .

Der Parameter a ist dabei ein Skalenparameter, die Form wird vom Parameter b festgelegt.

Bemerkung. Die hier gegebene Definition der Paretoverteilung wurde entnommen aus [Sch95, S. 35]. In der Literatur wird die Paretoverteilung nicht einheitlich definiert, die Verteilung mit der hier angegebenen Dichte wird gelegentlich auch als Nullpunkt-Paretoverteilung bezeichnet.

Bei der Paretoverteilung tritt der Fall auf, dass der Erwartungswert bzw. die Varianz nicht für jede beliebige Parameterwahl endlich sind. Das wird neben der Herleitung der Verteilungsfunktion in folgendem Satz gezeigt:

Satz 3.6. Sei die Zufallsvariable X paretoverteilt mit den Parametern $a > 0$ und $b > 0$. Dann gilt für den Erwartungswert $E(X)$, für die Varianz $V(X)$ und für die Verteilungsfunktion $F(x)$:

$$E(X) = \begin{cases} \frac{a}{b-1} & b > 1, \\ \infty & b \in (0, 1], \end{cases}$$

$$V(X) = \begin{cases} \infty & b \leq 2, \\ \frac{a^2 b}{(b-1)^2(b-2)} & b > 2, \end{cases}$$

$$F(x) = \begin{cases} 0 & x \leq 0, \\ 1 - \left(\frac{a}{a+x}\right)^b & x > 0. \end{cases}$$

Beweis.

Zuerst wird $E(X)$ berechnet, es gilt:

$$\begin{aligned} \int x \frac{b}{a} \left(\frac{a}{a+x}\right)^{b+1} dx &= a^b b \int \frac{x}{(a+x)^{b+1}} dx = \\ &\left(\text{partielle Integration: } u(x) = x, v'(x) = \frac{1}{(a+x)^{b+1}} \right) \\ &= a^b b \left(\frac{-x}{b(a+x)^b} + \int \frac{1}{b(a+x)^b} dx \right). \end{aligned}$$

Sei $b \neq 1$, dann gilt:

$$\begin{aligned} \int x \frac{b}{a} \left(\frac{a}{a+x}\right)^{b+1} dx &= a^b b \left(\frac{-x}{b(a+x)^b} + \frac{1}{b(1-b)(a+x)^{b-1}} \right) = \\ &= a^b \frac{-x(1-b) + a+x}{(a+x)^b(1-b)} = \\ &= a^b \frac{xb+a}{(a+x)^b(1-b)}. \end{aligned}$$

Daraus ergibt sich:

$$\begin{aligned} \int_0^\infty x \frac{b}{a} \left(\frac{a}{a+x}\right)^{b+1} dx &= a^b \frac{xb+a}{(a+x)^b(1-b)} \Big|_0^\infty = \\ &= \lim_{x \rightarrow \infty} a^b \frac{xb+a}{(a+x)^b(1-b)} - \underbrace{\lim_{x \rightarrow 0^+} a^b \frac{xb+a}{(a+x)^b(1-b)}}_{=\frac{a}{b-1}} = \\ \text{(Regel von de L'Hospital)} &= \lim_{x \rightarrow \infty} \underbrace{a^b \frac{b+a}{(b-1)(a+x)^{b-1}(1-b)}}_{=0} - \frac{a}{b-1} = \\ &= \begin{cases} \infty & b \in (0, 1), \\ 0 & b > 1. \end{cases} \end{aligned}$$

$$= \begin{cases} \infty & b \in (0, 1), \\ \frac{a}{b-1} & b > 1. \end{cases}$$

Sei $b = 1$, dann gilt:

$$\begin{aligned} \int x \frac{b}{a} \left(\frac{a}{a+x} \right)^{b+1} dx &= a \left(\frac{-x}{a+x} + \int \frac{1}{a+x} dx \right) = \\ &= \frac{-ax}{a+x} + a \cdot \log(a+x). \end{aligned}$$

Daraus ergibt sich:

$$\begin{aligned} \int_0^\infty x \frac{b}{a} \left(\frac{a}{a+x} \right)^{b+1} dx &= \left(\frac{-ax}{a+x} + a \cdot \log(a+x) \right) \Big|_0^\infty = \\ &= \underbrace{\lim_{x \rightarrow \infty} \left(\frac{-ax}{a+x} + a \cdot \log(a+x) \right)}_{=\infty} + \\ &\quad - \underbrace{\lim_{x \rightarrow 0^+} \left(\frac{-ax}{a+x} + a \cdot \log(a+x) \right)}_{=a \cdot \log(a)} = \\ &= \infty. \end{aligned}$$

Es gilt deshalb:

$$E(X) = \int_0^\infty x \frac{b}{a} \left(\frac{a}{a+x} \right)^{b+1} dx = \begin{cases} \frac{a}{b-1} & b > 1, \\ \infty & b \in (0, 1]. \end{cases}$$

Nun wird $E(X^2)$ berechnet:

$$\begin{aligned} \int x^2 \frac{b}{a} \left(\frac{a}{a+x} \right)^{b+1} dx &= a^b b \int \frac{x^2}{(a+x)^{b+1}} dx = \\ &\left(\text{partielle Integration: } u_1(x) = x^2, v_1'(x) = \frac{1}{(a+x)^{b+1}} \right) \\ &= a^b b \left(\frac{-x^2}{b(a+x)^b} + \int \frac{2x}{b(a+x)^b} dx \right) = \\ &= a^b \left(\frac{-x^2}{(a+x)^b} + \int \frac{2x}{(a+x)^b} dx \right). \end{aligned}$$

Sei $b = 1$, dann gilt:

$$\begin{aligned}
 \int x^2 \frac{b}{a} \left(\frac{a}{a+x} \right)^{b+1} dx &= a \left(\frac{-x^2}{a+x} + \int \underbrace{\frac{2x}{a+x}}_{=2 - \frac{2a}{x+a}} dx \right) = \\
 &= a \left(\frac{-x^2}{a+x} + \int \left(2 - \frac{2a}{x+a} \right) dx \right) = \\
 &= a \left(\frac{-x^2}{a+x} + 2x - 2a \cdot \log(x+a) \right).
 \end{aligned}$$

Setzt man die Integrationsgrenzen ein, dann gilt für $b = 1$:

$$\begin{aligned}
 \int_0^\infty x^2 \frac{b}{a} \left(\frac{a}{a+x} \right)^{b+1} dx &= a \left(\frac{-x^2}{a+x} + 2x - 2a \cdot \log(x+a) \right) \Big|_0^\infty = \\
 &= \underbrace{\lim_{x \rightarrow \infty} a \left(\frac{-x^2}{a+x} + 2x - 2a \cdot \log(x+a) \right)}_{=\infty} + \\
 &\quad - \underbrace{\lim_{x \rightarrow 0^+} a \left(\frac{-x^2}{a+x} + 2x - 2a \cdot \log(x+a) \right)}_{=-2a \cdot \log(a)} = \\
 &= \infty.
 \end{aligned}$$

Sei $b \neq 1$, dann gilt für $\int \frac{2x}{(a+x)^b} dx$:

$$\begin{aligned}
 \int \frac{2x}{(a+x)^b} dx &= \left(\text{partielle Integration: } u_2(x) = 2x, v_2'(x) = \frac{1}{(a+x)^b} \right) = \\
 &= \frac{2x}{(1-b)(a+x)^{b-1}} - \int \frac{2}{(1-b)(a+x)^{b-1}} dx = \\
 &= \begin{cases} \frac{2x}{(1-b)(a+x)^{b-1}} - \frac{2 \cdot \log(a+x)}{1-b} = \frac{-2x}{a+x} + 2 \cdot \log(a+x) & b = 2, \\ \frac{2x}{(1-b)(a+x)^{b-1}} - \frac{2}{(1-b)(2-b)(a+x)^{b-2}} & b \neq 2. \end{cases}
 \end{aligned}$$

Für $b = 2$ gilt dann:

$$\int_0^\infty x^2 \frac{b}{a} \left(\frac{a}{a+x} \right)^{b+1} dx = a^2 \left(\frac{-x^2}{(a+x)^2} - \frac{2x}{a+x} + 2 \cdot \log(a+x) \right) \Big|_0^\infty =$$

$$\begin{aligned}
&= \lim_{x \rightarrow \infty} \underbrace{a^2 \left(\frac{-x^2}{(a+x)^2} - \frac{2x}{a+x} + 2 \cdot \log(a+x) \right)}_{=\infty} + \\
&\quad - \lim_{x \rightarrow 0^+} \underbrace{a^2 \left(\frac{-x^2}{(a+x)^2} - \frac{2x}{a+x} + 2 \cdot \log(a+x) \right)}_{=2 \cdot \log(a)} = \\
&= \infty.
\end{aligned}$$

Für $b \notin \{1, 2\}$ gilt schließlich:

$$\begin{aligned}
&\int_0^\infty x^2 \frac{b}{a} \left(\frac{a}{a+x} \right)^{b+1} dx = \\
&= a^b \left(\frac{-x^2}{(a+x)^b} + \frac{2x}{(1-b)(a+x)^{b-1}} - \frac{2}{(1-b)(2-b)(a+x)^{b-2}} \right) \Big|_0^\infty = \\
&= \lim_{x \rightarrow \infty} \underbrace{a^b \left(\frac{-x^2}{(a+x)^b} + \frac{2x}{(1-b)(a+x)^{b-1}} - \frac{2}{(1-b)(2-b)(a+x)^{b-2}} \right)}_{= \begin{cases} 0 & b > 2, \\ \infty & b < 2. \end{cases}} + \\
&\quad - \lim_{x \rightarrow 0^+} \underbrace{a^b \left(\frac{-x^2}{(a+x)^b} + \frac{2x}{(1-b)(a+x)^{b-1}} - \frac{2}{(1-b)(2-b)(a+x)^{b-2}} \right)}_{=-\frac{2a^b}{(1-b)(2-b)a^{b-2}} = -\frac{2a^2}{(1-b)(2-b)}} = \\
&= \begin{cases} \frac{2a^2}{(1-b)(2-b)} & b > 2, \\ \infty & b < 2. \end{cases}
\end{aligned}$$

Daraus ergibt sich für $E(X^2)$:

$$E(X^2) = \int_0^\infty x^2 \frac{b}{a} \left(\frac{a}{a+x} \right)^{b+1} dx = \begin{cases} \frac{2a^2}{(1-b)(2-b)} & b > 2, \\ \infty & b \in (0, 2]. \end{cases}$$

Für $V(X)$ gilt dann also:

$$V(X) = E(X^2) - (E(X))^2 = \begin{cases} \infty & b \leq 2, \\ \frac{2a^2}{(1-b)(2-b)} - \left(\frac{a}{b-1} \right)^2 = \frac{a^2 b}{(b-1)^2(b-2)} & b > 2. \end{cases}$$

Für die Verteilungsfunktion $F(t)$ gilt für $t > 0$:

$$\begin{aligned} F(t) &= \int_0^t \frac{b}{a} \left(\frac{a}{a+x} \right)^{b+1} dx = \frac{-a^b}{(a+x)^b} \Big|_0^t = \\ &= -\frac{a^b}{(a+x)^b} + \frac{a^b}{a^b} = \\ &= 1 - \left(\frac{a}{a+t} \right)^b. \end{aligned}$$

Für $t \leq 0$ gilt:

$$F(t) = 0.$$

□

3.1.5 Verteilungen zur Modellierung von Großschäden

Großschäden stellen Versicherungsunternehmen vor Probleme, denn diese hohen Schäden treten nur mit einer geringen Wahrscheinlichkeit auf, aber im Falle eines Auftretens muss das Versicherungsunternehmen große Geldbeträge für die Schadensregulierung aufwenden. Der Begriff des Großschadens kann nicht genau definiert bzw. anhand eines bestimmten Geldbetrages festgelegt werden, sondern bezieht sich immer auf einen konkreten Bestand von Risiken. Versicherungsunternehmen sprechen dann von Großschäden, wenn die Höhe eines Schadens im Vergleich zu den restlichen Schäden des Bestandes außergewöhnlich hoch ist (vgl. [FW01, S. 299f]).

Bei der Modellierung von Risiken muss darauf geachtet werden, dass die Möglichkeit des Auftretens eines Großschadens ausreichend berücksichtigt wird. Sind hohe Schäden möglich, dann muss für die Modellierung eine Verteilung verwendet werden, die Großschäden nicht unterschätzt. Das heißt, dass die Dichtefunktion der Verteilung für große Werte nur langsam gegen 0 konvergieren darf bzw. dass die Verteilungsfunktion nur langsam gegen 1 konvergieren darf. Mathematisch exakt spricht man bei einer Verteilung, die diese Eigenschaft besitzt, von einer heavy-tailed Verteilung, dieser Begriff wird im Folgenden exakt definiert.

Ein Beispiel für einen Versicherungszweig, in dem Großschäden auftreten können ist die Feuerversicherung (vgl. [FW01, S. 299]).

Definition 3.9.

Sei $F(x)$ die Verteilungsfunktion einer Wahrscheinlichkeitsverteilung. Die Verteilung

heißt heavy-tailed, wenn folgendes gilt:

$$\lim_{x \rightarrow \infty} e^{kx}(1 - F(x)) = \infty \quad \forall k > 0.$$

Bemerkung. In der Literatur werden Heavy-tailed-Verteilungen gelegentlich auch anders als hier definiert, dabei sind die angegebenen Definitionen nicht notwendigerweise äquivalent zur hier angeführten Definition, die aus [Asm03] entnommen wurde.

Im Folgenden werden die Paretoverteilung und die Exponentialverteilung auf die Heavy-tailed-Eigenschaft untersucht:

Satz 3.7. *Die Paretoverteilung ist heavy-tailed.*

Beweis. nach [6, S. 28]

Sei $F(x)$ die Verteilungsfunktion der Paretoverteilung, dann gilt für alle $k \in \mathbb{R}, k > 0$:

$$\begin{aligned} \lim_{x \rightarrow \infty} e^{kx}(1 - F(x)) &= \lim_{x \rightarrow \infty} e^{kx} \left(1 - \left(1 - \left(\frac{a}{a+x} \right)^b \right) \right) = \\ &= a^b \lim_{x \rightarrow \infty} \frac{e^{kx}}{(a+x)^b}. \end{aligned}$$

Um diesen Ausdruck zu berechnen wird folgende Abschätzung verwendet: Ist $a+x > 1$ dann gilt:

$$\frac{1}{(a+x)^b} \geq \frac{1}{(a+x)^c} \quad \forall c : c \geq b > 0.$$

Sei $n \in \mathbb{N}$ so, dass $n < b \leq n+1$, dann gilt deshalb:

$$\lim_{x \rightarrow \infty} \frac{e^{kx}}{(a+x)^b} \geq \lim_{x \rightarrow \infty} \frac{e^{kx}}{(a+x)^{n+1}}.$$

Der rechte Grenzwert lässt sich berechnen indem man die Regel von de L'Hospital $n+1$ mal anwendet, es ergibt sich dann:

$$\lim_{x \rightarrow \infty} \frac{e^{kx}}{(a+x)^{n+1}} = \lim_{x \rightarrow \infty} \frac{k^{n+1} e^{kx}}{(n+1)!} = \infty \quad \forall k > 0.$$

Insgesamt ergibt sich daraus:

$$\lim_{x \rightarrow \infty} \frac{e^{kx}}{(a+x)^b} = \infty \quad \forall k > 0.$$

Damit ist bewiesen, dass die Paretoverteilung heavy-tailed ist. □

Satz 3.8. *Die Exponentialverteilung ist nicht heavy-tailed.*

Beweis.

Sei $F(x)$ die Verteilungsfunktion der Exponentialverteilung, dann gilt:

$$\begin{aligned} \lim_{x \rightarrow \infty} e^{kx}(1 - F(x)) &= \lim_{x \rightarrow \infty} e^{kx}(1 - (1 - e^{-\lambda x})) = \\ &= \lim_{x \rightarrow \infty} e^{k-\lambda x} = \\ &= \begin{cases} 0 & k < \lambda, \\ 1 & k = \lambda, \\ \infty & k > \lambda. \end{cases} \end{aligned}$$

Die Exponentialverteilung ist also nicht heavy-tailed. $\square$

Es wurde gezeigt, dass die Paretoverteilung heavy-tailed ist, die Exponentialverteilung hingegen nicht. Das heißt für die Modellierung von Risiken, dass sich die Paretoverteilung für Risiken, die Großschäden verursachen können, besser eignet als die Exponentialverteilung, da bei der Paretoverteilung hohe Schäden anhand der Verteilungsfunktion wahrscheinlicher sind als bei der Exponentialverteilung. Wie in [13, S. 30] angeführt wird sind die Gammaverteilung und die Inverse Normalverteilung wie die Exponentialverteilung nicht heavy-tailed, diese Verteilungen eignen sich deshalb nur zur Modellierung von Risiken mit kleinen und mittleren Schadenshöhen.

3.1.6 Ein Modell für den Gesamtschaden mit der Logarithmischen Normalverteilung

Die Gammaverteilung und die Inverse Normalverteilung besitzen die Eigenschaft, dass sie invariant unter Faltungen sind. Das heißt für versicherungsmathematische Anwendungen, dass der Gesamtschaden eines Bestandes von gamma- oder invers-normalverteilten Risiken wiederum gamma- bzw. invers-normalverteilt ist. Allerdings besitzen nur wenige Wahrscheinlichkeitsverteilungen diese Eigenschaft. Da die Berechnung von Faltungen mitunter schwierig und rechenaufwändig ist wird in [Mac02] und [13] ein Modell vorgestellt, in dem nicht die einzelnen Risiken mittels Wahrscheinlichkeitsverteilungen beschrieben werden und die Verteilung des Gesamtschadens über die Faltung der Einzelschadenverteilungen bestimmt wird, sondern in dem nur der Gesamtschaden mit einer geeigneten Verteilung beschrieben wird. Als Verteilung wird dabei die Logarithmische

Normalverteilung verwendet, dies wird damit begründet, dass die Logarithmische Normalverteilung eine ähnliche Form wie die Gammaverteilung und die Inverse Normalverteilung besitzt und dass die Logarithmische Normalverteilung rechentechnische Vorteile aufweist, da sich manche Rechnungen auf die Normalverteilung zurückführen lassen (vgl. [13, S. 16]).

Die Logarithmische Normalverteilung ist folgendermaßen definiert:

Definition 3.10.

Als Logarithmische Normalverteilung $LN(\mu, \sigma^2)$ mit den Parametern $\mu \in \mathbb{R}$ und $\sigma^2 \in \mathbb{R}$ wird die Wahrscheinlichkeitsverteilung mit der Dichte

$$f(x) = \begin{cases} \frac{1}{\sqrt{2\pi\sigma^2}x} e^{-\frac{(\log(x)-\mu)^2}{2\sigma^2}} & x > 0, \\ 0 & x \leq 0, \end{cases}$$

bezeichnet. Der Parameter σ^2 muss $\sigma^2 > 0$ erfüllen.

Der Parameter μ ist hier ein Skalenparameter, σ^2 ist ein Formparameter.

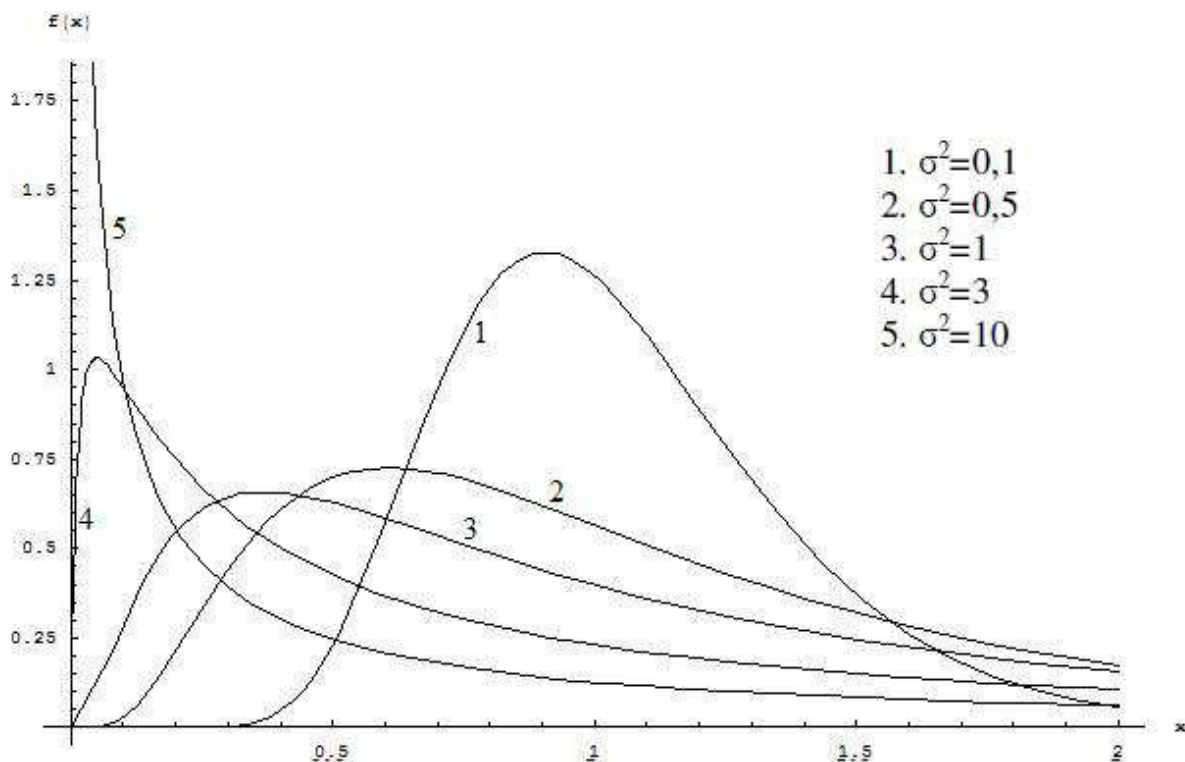


Abbildung 3.5: Die Dichtefunktion der Logarithmischen Normalverteilung mit $\mu = 0$ und verschiedenen Werten für σ^2 .

Für die Logarithmische Normalverteilung gilt folgende Äquivalenz, die das Rechnen mit dieser Verteilung erleichtert (und worauf auch ihr Name zurückzuführen ist): Die Zufallsvariable Y ist genau dann normalverteilt mit den Parametern μ und σ^2 wenn die Zufallsvariable $X := e^Y$ logarithmisch normalverteilt mit den Parametern μ und σ^2 ist. Diese Äquivalenz kann benutzt werden, um den Erwartungswert und die Varianz der Logarithmischen Normalverteilung herzuleiten, was im Beweis des folgenden Satzes benutzt wird:

Satz 3.9. *Sei die Zufallsvariable X logarithmisch normalverteilt mit den Parametern μ und σ^2 , dann gilt für den Erwartungswert, die Varianz und die Verteilungsfunktion:*

$$\begin{aligned} E(X) &= e^{\mu + \frac{\sigma^2}{2}}, \\ V(X) &= e^{2\mu + \sigma^2} (e^{\sigma^2} - 1), \\ F(x) &= \begin{cases} 0 & x \leq 0, \\ \Phi\left(\frac{\log(x) - \mu}{\sigma}\right) & x > 0. \end{cases} \end{aligned}$$

(Φ bezeichnet dabei die Verteilungsfunktion der Standardnormalverteilung.)

Beweis. (nach [AW05, S. 76] und [Sch95, S. 45])

Sei die Zufallsvariable Y normalverteilt mit den Parametern μ und σ^2 , dann gilt für den Erwartungswert der logarithmisch normalverteilten Zufallsvariable X :

$$\begin{aligned} E(X) &= E(e^Y) = \\ &= \int_{-\infty}^{\infty} e^y \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(y-\mu)^2}{2\sigma^2}} dy = \\ &\quad \left(\text{Substitution: } u = \frac{y - \mu}{\sigma} \Leftrightarrow y = u\sigma + \mu \right) \\ &= \int_{-\infty}^{\infty} e^{u\sigma + \mu} \frac{1}{\sqrt{2\pi}} e^{-\frac{u^2}{2}} du = \\ &= e^{\mu} \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} \underbrace{e^{-\frac{u^2}{2} + u\sigma}}_{= e^{-\frac{1}{2}(u-\sigma)^2 + \frac{\sigma^2}{2}}} du = \\ &= e^{\mu} \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(u-\sigma)^2 + \frac{\sigma^2}{2}} du = \end{aligned}$$

$$\begin{aligned}
&= e^{\mu + \frac{\sigma^2}{2}} \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(u-\sigma)^2} du = \\
&\quad (\text{Substitution: } t = u - \sigma) \\
&= e^{\mu + \frac{\sigma^2}{2}} \underbrace{\int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}t^2} dt}_{=1} = \\
&= e^{\mu + \frac{\sigma^2}{2}}.
\end{aligned}$$

Um $E(X^2)$ zu berechnen müssen nur die Rechenregeln für Potenzen angewendet werden, es gilt:

$$E(X^2) = E((e^Y)^2) = E(e^{2Y}).$$

Da die Zufallsvariable Y normalverteilt mit den Parametern μ und σ^2 ist, ist die Zufallsvariable $2Y$ aufgrund der Eigenschaften der Normalverteilung ebenfalls normalverteilt mit den Parametern 2μ und $4\sigma^2$. Es gilt daher:

$$E(X^2) = e^{2\mu + \frac{4\sigma^2}{2}} = e^{2\mu + 2\sigma^2}.$$

Für die Varianz gilt deshalb:

$$\begin{aligned}
V(X) &= E(X^2) - (E(X))^2 = \\
&= e^{2\mu + 2\sigma^2} - \left(e^{\mu + \frac{\sigma^2}{2}}\right)^2 = \\
&= e^{2\mu + 2\sigma^2} - e^{2\mu + \sigma^2} = \\
&= e^{2\mu + \sigma^2} (e^{\sigma^2} - 1).
\end{aligned}$$

Sei $x > 0$, dann gilt für die Verteilungsfunktion der Logarithmischen Normalverteilung:

$$F(x) = P(X < x) = P(e^Y < x) = P(Y < \log(x)).$$

Das heißt, die Verteilungsfunktion $F(x)$ der Logarithmischen Normalverteilung kann mit Hilfe der Verteilungsfunktion der Normalverteilung ausgedrückt werden, es gilt:

$$F(x) = P(Y < \log(x)) = \frac{1}{\sqrt{2\pi\sigma^2}} \int_{-\infty}^{\log(x)} e^{-\frac{(t-\mu)^2}{2\sigma^2}} dt.$$

Ausgedrückt mit der Verteilungsfunktion der Standardnormalverteilung heißt das:

$$F(x) = \Phi\left(\frac{\log(x) - \mu}{\sigma}\right).$$

Sei $x \leq 0$, dann gilt für die Verteilungsfunktion der Logarithmischen Normalverteilung:

$$F(x) = 0.$$

Insgesamt gilt also:

$$F(x) = \begin{cases} 0 & x \leq 0, \\ \Phi\left(\frac{\log(x) - \mu}{\sigma}\right) & x > 0. \end{cases}$$

□

Die Logarithmische Normalverteilung wird nun benutzt um den Gesamtschaden eines Bestandes zu beschreiben. Um den Einfluss der Gesamtversicherungssumme v in die Modellierung miteinzubeziehen wird allerdings nicht der Gesamtschaden S mit der Logarithmischen Normalverteilung beschrieben, sondern die Zufallsvariable Z , die als $Z := \frac{S}{v}$ definiert wird. (Unter der Versicherungssumme wird dabei jener Betrag verstanden, der im Schadensfall maximal ausgezahlt wird, auch wenn ein aufgetretener Schaden diesen Betrag übersteigt.) Hier wird Z anstatt dem Gesamtschaden S verwendet, weil zur Modellierung Daten aus vergangenen Perioden herangezogen werden müssen. Da sich die Anzahl an Versicherungsverträgen aber von Periode zu Periode ändern kann, was zur Folge hat, dass sich auch der Gesamtschaden ändert, wird hier ein Modell für eine auf das Volumen bezogene Zufallsvariable erstellt. (Wie im nächsten Abschnitt beschrieben wird, kann dieses Vorgehen auch bei der Erstellung von Modellen für nicht homogene Bestände angewendet werden.)

Außerdem wird die Zufallsvariable T eingeführt, die normalverteilt mit den Parametern μ und $\frac{\sigma^2}{v}$ ist. Es gilt dann: $Z = e^T$. Deswegen gilt für den Erwartungswert und für die Varianz von Z (vgl. [Mac02, S. 56]):

$$E(Z) = e^{\mu + \frac{\sigma^2}{2v}},$$

$$V(Z) = e^{2\mu + \frac{\sigma^2}{v}} \left(e^{\frac{\sigma^2}{v}} - 1 \right).$$

Um die Parameter μ und σ^2 zu schätzen wird auf die Daten vergangener Perioden zurückgegriffen. Seien $v_1, v_2, \dots, v_j$ die Gesamtversicherungssummen in j betrachteten

Perioden und sei z_i die (inflationsbereinigte) Realisierung der Zufallsvariable Z in der i . Periode ($i \in \{1, 2, \dots, j\}$). Dann sind $\hat{\mu}$ und $\hat{\sigma}^2$ erwartungstreue Schätzer für μ bzw. σ^2 (vgl. [Mac02, S. 56]):

$$\hat{\mu} = \frac{1}{\sum_{i=1}^j v_i} \sum_{i=1}^j v_i \log(z_i),$$

$$\hat{\sigma}^2 = \frac{1}{j-1} \sum_{i=1}^j v_i (\log(z_i) - \hat{\mu})^2.$$

Dieses Modell erlaubt zwar einfache Berechnungen (besonders durch die Rückführung auf die Normalverteilung), allerdings ist eine umfangreiche Datenbasis notwendig, um es anzuwenden.

3.1.7 Inhomogene Bestände

Vernünftig einsetzen lässt sich das individuelle Modell nur in (weitgehend) homogenen Beständen. Gerade in der Schadenversicherung sind die Bestände aber üblicherweise sehr inhomogen. Eine Möglichkeit wäre es, die Risiken so in Teilbestände aufzuteilen, dass diese einerseits die Bedingung der Unabhängigkeit, andererseits die Bedingung der gleichen Verteilung erfüllen. Die Gesamtschadenverteilung würde man über die Faltung der Verteilungen der einzelnen Bestände erhalten. Dies wäre möglicherweise nicht mehr analytisch, sondern nur mehr numerisch durchführbar und rechenaufwändig. Außerdem hätte man eventuell das Problem, dass die einzelnen Bestände wegen der geforderten Homogenität sehr klein wären und deshalb kein der Realität entsprechendes Modell beschreiben würden. Würde in einem Bestand ein sehr hoher Schaden auftreten, dann könnte er nicht mehr durch die anderen Risiken dieses Bestandes kompensiert werden (vgl. [Mac02, S. 77]).

Ein Modell für inhomogene Bestände, in denen die Risiken unabhängig sind und verschieden hohe Versicherungssummen besitzen, wird in [13, S. 5f] beschrieben:

Der Ausgangspunkt ist hier ein sogenanntes „Referenzrisiko“ X_0 , das eine Versicherungssumme u_0 besitzt. Weiters werden m und s^2 folgendermaßen definiert:

$$m := \frac{1}{u_0} \cdot E(X_0),$$

$$s^2 := \frac{1}{u_0} \cdot V(X_0).$$

Wird nun ein Risiko X_k betrachtet, dessen Versicherungssumme u_k von u_0 abweicht,

dann wird angenommen, dass sich der Erwartungswert und die Varianz des Risikos X_k im Verhältnis $\frac{u_k}{u_0}$ zum Erwartungswert und zur Varianz von X_0 verhalten, also:

$$E(X_k) = \frac{u_k}{u_0} \cdot E(X_0) = m \cdot u_k \quad \forall k = 1, 2, \dots, n;$$

$$V(X_k) = \frac{u_k}{u_0} \cdot V(X_0) = s^2 \cdot u_k \quad \forall k = 1, 2, \dots, n.$$

Die Versicherungssumme des gesamten Bestandes v ergibt sich folgendermaßen:

$$v := \sum_{k=1}^n u_k.$$

Deshalb gilt:

$$E(S) = v \cdot m,$$

$$V(S) = v \cdot s^2.$$

Schätzen der Parameter

Wie für homogene Bestände besprochen wird auch im oben beschriebenen Modell eine auf die Bestandsgröße bezogene Größe eingeführt. Dies ist der Schadensatz Z , der folgendermaßen definiert ist (vgl. [Mac02, S. 41]):

$$Z := \frac{S}{v} = \frac{\sum_{k=1}^n X_k}{\sum_{k=1}^n u_k}.$$

Analog zum Schadenbedarf im homogenen Fall gilt:

$$E(Z) = m.$$

Der Schadensatz ist auch hier ein erwartungstreuer Schätzer von m . Für die Varianz des Schadensatzes ergibt sich:

$$V(Z) = \frac{s^2}{v}.$$

Wie im homogenen Fall zeigt sich auch hier ein Ausgleichseffekt, je größer die Gesamtversicherungssumme v ist, desto kleiner ist die Varianz des Schadensatzes.

Bei bekannten Realisierungen $x_1, x_2, \dots, x_n$ der Risiken $X_1, X_2, \dots, X_n$ sind wie im ho-

mogenen Fall $\hat{m}$ und $\hat{s}^2$ erwartungstreue Schätzer für m und s^2 (vgl. [Mac02, S. 41]):

$$\hat{m} = \frac{1}{v} \sum_{i=1}^n x_i,$$

$$\hat{s}^2 = \frac{1}{n-1} \sum_{i=1}^n u_i \left(\frac{x_i}{u_i} - \hat{m} \right)^2.$$

Sind die Realisierungen des Schadensatzes und die Bestandsgrößen mehrerer Jahre bekannt, dann können mit diesen Daten Schätzungen vorgenommen werden: Sei j die Anzahl der betrachteten Perioden, z_i die (inflationsbereinigte) Realisierung des Schadensatzes in der i . Periode ($i \in \{1, 2, \dots, j\}$) und v_i die (ebenfalls inflationsbereinigte) Versicherungssumme in dieser Periode. Dann sind $\hat{m}$ und $\hat{s}^2$ erwartungstreue Schätzer für m bzw. s^2 (vgl. [Mac02, S. 42]):

$$\hat{m} = \frac{1}{\sum_{i=1}^j v_i} \sum_{i=1}^j v_i z_i,$$

$$\hat{s}^2 = \frac{1}{j-1} \sum_{i=1}^j v_i (z_i - \hat{m})^2.$$

3.2 Kollektives Modell

3.2.1 Allgemeines zum Modell

Anders als im individuellen Modell werden im kollektiven Modell nur tatsächliche Schäden erfasst. Ausgangspunkt sind hier nicht mehr die einzelnen Versicherungsverträge, sondern die Zahl der Schäden in einer Betrachtungsperiode. Es wird in diesem Modell nicht erhoben, welche Verträge Schäden verursacht haben und welche nicht. Im Gegensatz zum individuellen Modell ist deshalb nicht im Vorhinein klar, aus wie vielen Summanden sich der Gesamtschaden zusammensetzt, da die Zahl der Schäden vom Zufall abhängt (vgl. [Sch95, S. 108]).

Definition 3.11.

Als Schadenzahl wird eine Zufallsvariable N bezeichnet, die die Anzahl der Schadensfälle in einer Betrachtungsperiode angibt.

Die Zufallsvariable S sei der Gesamtschaden in dieser Periode.

Die Zufallsvariablen $X_1, X_2, \dots, X_N$ geben in diesem Modell die Höhe der einzelnen Schäden an. Der Gesamtschaden ergibt sich deshalb aus der Summe über alle X_i :

$$S = \sum_{i=1}^N X_i.$$

Für $N = 0$ setzt man:

$$S := 0.$$

In diesem Modell geht man, falls nicht anders angegeben, von folgenden Annahmen aus:

- Die Zufallsvariablen $N, X_1, X_2, \dots$ sind unabhängig.
- Die Zufallsvariablen $X_1, X_2, \dots$ sind identisch verteilt.

Die Unabhängigkeit der Schadenshöhen $X_1, X_2, \dots$ von der Schadenzahl N kann für die meisten Versicherungssparten als realistisch betrachtet werden, allerdings sind auch Ausnahmen denkbar: Beispielsweise kann es bei der Kfz-Haftpflichtversicherung im Winter bei eisigen Fahrbahnen zu vielen kleinen Schäden kommen (vgl. [Mac02, S. 109]).

Bei der Unabhängigkeit der Schadenshöhen $X_1, X_2, \dots$ tritt in diesem Modell kein Problem auf, da Schäden, die einander bedingen, zu einem Schadensfall zusammengefasst werden können (vgl. [Wol88, S. 52]).

Sei p_n die Wahrscheinlichkeit, dass in der Betrachtungsperiode genau n Schäden eintreten, also:

$$p_n := P(N = n).$$

Weiters sei F_i die Verteilungsfunktion der Zufallsvariable X_i . Da die Zufallsvariablen $X_1, X_2, \dots$ laut Annahme identisch verteilt sind kann die Verteilungsfunktion dieser Zufallsvariable mit F bezeichnet werden: $F := F_1 (= F_2 = F_3 = \dots)$. Außerdem soll zur Vereinfachung die typische Schadenshöhe X eingeführt werden, die die gleiche Verteilung wie die Zufallsvariablen $X_1, X_2, \dots$ haben soll.

Sei F_S die Verteilungsfunktion des Gesamtschadens S :

$$F_S(x) = P(S \leq x).$$

Aufgrund der Formel der totalen Wahrscheinlichkeit gilt für $P(S \leq x)$:

$$\begin{aligned} P(S \leq x) &= \sum_{n=0}^{\infty} P(S \leq x | N = n) \cdot P(N = n) = \\ &= \sum_{n=0}^{\infty} P\left(\sum_{i=1}^n X_i \leq x\right) \cdot p_n. \end{aligned}$$

Da alle Zufallsvariablen X_i identisch verteilt sind kann die Verteilungsfunktion von $\sum_{i=1}^n X_i$ als n -fache Faltung der Verteilungsfunktion F geschrieben werden:

$$P\left(\sum_{i=1}^n X_i \leq x\right) = F^{*n}(x).$$

Deshalb gilt:

$$F_S(x) = \sum_{n=0}^{\infty} F^{*n}(x) \cdot p_n.$$

(vgl. [Sch95, S. 80])

Dabei soll folgende Konvention gelten:

$$F^{*0}(x) := \begin{cases} 0 & x < 0, \\ 1 & x \geq 0. \end{cases}$$

Sei $f(x)$ die Dichte der Verteilung der typischen Schadenshöhe X , dann gilt für die

Dichte des Gesamtschadens $f_S(x)$:

$$f_S(x) = \sum_{n=0}^{\infty} f^{*n}(x) \cdot p_n.$$

Es gilt hier:

$$f^{*0}(x) := \begin{cases} 0 & x \neq 0, \\ 1 & x = 0. \end{cases}$$

Da in diesem Modell verglichen mit dem individuellen Modell nicht a priori klar ist, aus wie vielen Einzelschäden sich der Gesamtschaden zusammensetzt, wird die Zufallsvariable I_n eingeführt, die als Indikator für die Schadenzahl dient und zur Berechnung des Erwartungswertes und der Varianz des Gesamtschadens benötigt wird:

$$I_n = \begin{cases} 1 & \text{falls } N = n, \\ 0 & \text{sonst.} \end{cases}$$

Für I_n gilt dann:

$$E(I_n) = p_n.$$

Nun können der Erwartungswert und die Varianz des Gesamtschadens berechnet werden:

Satz 3.10. *Für den Erwartungswert und die Varianz des Gesamtschadens im kollektiven Modell gilt:*

$$E(S) = E(N)E(X),$$

$$V(S) = E(N)V(X) + V(N)(E(X))^2.$$

Beweis. nach [Sch02, S. 166f]

Für $E(S)$ gilt:

$$\begin{aligned} E(S) &= E\left(\sum_{i=1}^N X_i\right) = \\ &= E\left(\sum_{n=1}^{\infty} I_n \sum_{i=1}^n X_i\right) = \\ &= \underbrace{E\left(\sum_{n=1}^{\infty} I_n\right)}_{= \sum_{n=1}^{\infty} E(I_n)} \underbrace{E\left(\sum_{i=1}^n X_i\right)}_{= nE(X)} = \end{aligned}$$

$$\begin{aligned}
&= \sum_{n=1}^{\infty} \underbrace{E(I_n)}_{=p_n} n E(X) = \\
&= \sum_{n=1}^{\infty} \underbrace{p_n n}_{=E(N)} E(X) = \\
&= E(N) E(X).
\end{aligned}$$

(Bei der Berechnung von $E(S)$ wurde verwendet, dass für unabhängige Zufallsvariablen der Erwartungswert des Produkts gleich dem Produkt der Erwartungswerte ist, also dass $E(YZ) = E(Y)E(Z)$ für unabhängige Zufallsvariablen Y, Z gilt)

Für $E(S^2)$ gilt:

$$\begin{aligned}
E(S^2) &= E \left(\left(\sum_{i=1}^N X_i \right)^2 \right) = \\
&= E \left(\sum_{n=1}^{\infty} I_n \left(\sum_{i=1}^n X_i \right)^2 \right) = \\
&= \underbrace{E \left(\sum_{n=1}^{\infty} I_n \right)}_{=\sum_{n=1}^{\infty} p_n} \underbrace{E \left(\left(\sum_{i=1}^n X_i \right)^2 \right)}_{=V(\sum_{i=1}^n X_i) + (E(\sum_{i=1}^n X_i))^2} = \\
&= \sum_{n=1}^{\infty} p_n \left(\underbrace{V \left(\sum_{i=1}^n X_i \right)}_{=nV(X)} + \left(\underbrace{E \left(\sum_{i=1}^n X_i \right)}_{=nE(X)} \right)^2 \right) = \\
&= \sum_{n=1}^{\infty} p_n (nV(X) + n^2(E(X))^2) = \\
&= \underbrace{\sum_{n=1}^{\infty} p_n n V(X)}_{=E(N)V(X)} + \underbrace{\sum_{n=1}^{\infty} p_n n^2 (E(X))^2}_{=E(N^2)(E(X))^2} = \\
&= E(N)V(X) + E(N^2)(E(X))^2.
\end{aligned}$$

(Auch hier wurde verwendet, dass für unabhängige Zufallsvariablen der Erwartungswert des Produkts gleich dem Produkt der Erwartungswerte ist.)

Daraus ergibt sich für $V(S)$:

$$\begin{aligned}
 V(S) &= E(S^2) - (E(S))^2 = \\
 &= E(N)V(X) + E(N^2)(E(X))^2 - (E(N)E(X))^2 = \\
 &= E(N)V(X) + E(N^2)(E(X))^2 - (E(N))^2(E(X))^2 = \\
 &= E(N)V(X) + (E(X))^2 \underbrace{(E(N^2) - (E(N))^2)}_{=V(N)} = \\
 &= E(N)V(X) + V(N)(E(X))^2.
 \end{aligned}$$

□

3.2.2 Verteilungen der Schadenzahl

Im kollektiven Modell ist es weniger wichtig, wie groß die Wahrscheinlichkeit ist, dass einzelne Risiken Schäden verursachen, es ist vielmehr von Bedeutung, wie die Anzahl der Schäden des Bestandes verteilt sind (vgl. [Wol88, S. 70]).

In diesem Abschnitt werden die wichtigsten Verteilungen der Schadenzahl N und ihre Eigenschaften behandelt.

Binomialverteilung

Die Binomialverteilung ist folgendermaßen definiert:

Definition 3.12.

Eine diskrete Zufallsvariable N heißt binomialverteilt mit den Parametern n und p ($B(n, p)$ -verteilt), wenn

$$P(N = k) = \binom{n}{k} p^k (1 - p)^{n-k} \quad k = 0, 1, 2, \dots$$

Der Parameter n muss $n \in \mathbb{N}$ erfüllen, für den Parameter p muss $p \in (0, 1)$ gelten.

Für den Erwartungswert und die Varianz einer $B(n, p)$ -verteilten Zufallsvariable X gilt:

$$E(N) = np,$$

$$V(N) = np(1 - p).$$

Wird die Binomialverteilung zur Modellierung der Schadenzahl N verwendet, dann entspricht der Parameter n der Maximalanzahl an Schäden im betrachteten Intervall,

deshalb kann die Binomialverteilung nur dann eingesetzt werden, wenn die Anzahl der auftretenden Schäden nach oben begrenzt ist. Geeignet ist die Binomialverteilung vor allem dann, wenn die Schadenzahl für kleine homogene Betände modelliert werden soll (vgl. [13, S. 19f]). Da der Parameter n die Maximalanzahl an Schäden angibt kann nur der Parameter p zum Anpassen an beobachtete Daten verwendet werden, weshalb die Binomialverteilung nicht so flexibel an beobachtete Daten angepasst werden kann wie eine Verteilung mit zwei Parametern. In [Wol88, S. 73] wird diese Verteilung vor allem dann empfohlen, wenn die erwartete Schadenzahl größer als deren Varianz ist.

Poissonverteilung

Eine bedeutende und in der Praxis vielfach eingesetzte Verteilung der Schadenzahl ist die Poissonverteilung. Sie ist folgendermaßen definiert:

Definition 3.13.

Eine diskrete Zufallsvariable N heißt poissonverteilt mit dem Parameter $\lambda > 0$, wenn

$$P(N = k) = \frac{\lambda^k}{k!} e^{-\lambda} \quad k = 0, 1, 2, \dots$$

Der Parameter λ wird als Intensität bezeichnet.

Sei N poissonverteilt mit dem Parameter λ , dann gilt:

$$E(N) = \lambda,$$

$$V(N) = \lambda.$$

In [Wol88, S. 71f] werden folgende drei Voraussetzungen für das Auftreten der Poissonverteilung genannt:

1. Die Wahrscheinlichkeit, dass in einem kurzen Intervall der Länge Δ genau ein Schaden eintritt ist ungefähr gleich $\lambda\Delta$:

$$P(\text{Anzahl der Schäden im Zeitintervall der Länge } \Delta = 1) = \lambda\Delta + \mathcal{O}(\Delta).$$

2. Die Wahrscheinlichkeit, dass mehr als ein Schaden im kurzen Intervall der Länge Δ eintritt, ist klein gegen die Wahrscheinlichkeit, dass genau ein Schaden eintritt:

$$P(\text{Anzahl der Schäden im Zeitintervall der Länge } \Delta > 1) = \mathcal{O}(\Delta).$$

3. Die Schadenzahlen in disjunkten Zeitintervallen sind unabhängig.

Satz 3.11. *Wenn die Bedingungen 1–3 gelten, dann ist die Schadenzahl in einem Intervall der Länge t poissonverteilt mit dem Parameter $\nu = \lambda t$. Es gilt dann für die Schadenzahl N_t im Intervall $[0, t]$ und für ein $k \in \mathbb{N}$, k beliebig:*

$$P(N_t = k) = \frac{(\lambda t)^k}{k!} \cdot e^{-\lambda t}$$

Beweis. (nach [Wol88, S. 72])

Zuerst wird das Intervall $[0, t]$ in gleich lange, disjunkte Teilintervalle aufgeteilt, ihre Länge beträgt:

$$\Delta = \frac{t}{n} \quad n \in \mathbb{N} \text{ beliebig.}$$

Ist n groß (d. h. die Länge Δ des Intervalls ist klein), dann fällt in jedes Teilintervall kein Schaden oder genau ein Schaden, der Fehler beträgt $\mathcal{O}(\Delta)$. Nach der Bedingung 1 ist die Wahrscheinlichkeit, dass ein Schaden in einem Intervall auftritt $\lambda\Delta$.

Aufgrund der Bedingung 3 gilt für $k \leq n$: Die Wahrscheinlichkeit von genau k Schadensfällen im Intervall $[0, t]$ ist:

$$P(N_t = k) = \binom{n}{k} (\lambda\Delta)^k (1 - \lambda\Delta)^{n-k} = \binom{n}{k} \left(\frac{\lambda t}{n}\right)^k \left(1 - \frac{\lambda t}{n}\right)^{n-k}.$$

Für $n \rightarrow \infty$ gilt deswegen:

$$\begin{aligned} P(N_t = k) &= \lim_{n \rightarrow \infty} \binom{n}{k} \left(\frac{\lambda t}{n}\right)^k \left(1 - \frac{\lambda t}{n}\right)^{n-k} = \\ &= \lim_{n \rightarrow \infty} \frac{n \cdot (n-1) \cdot \dots \cdot (n-k+1)}{k!} \left(\frac{\lambda t}{n}\right)^k \left(1 - \frac{\lambda t}{n}\right)^{n-k} \\ &= \lim_{n \rightarrow \infty} = \underbrace{\frac{1}{k!}}_{\rightarrow \frac{1}{k!}} \underbrace{\frac{n \cdot (n-1) \cdot \dots \cdot (n-k+1)}{n^k}}_{\rightarrow 1} \underbrace{(\lambda t)^k}_{\rightarrow (\lambda t)^k} \underbrace{\left(1 - \frac{\lambda t}{n}\right)^{n-k}}_{\rightarrow e^{-\lambda t}} \\ &= \frac{(\lambda t)^k}{k!} \cdot e^{-\lambda t}. \end{aligned}$$

□

Wie bei der Binomialverteilung steht auch bei der Poissonverteilung nur ein Parameter zur Modellierung der Schadenzahl zur Verfügung, wodurch die Anpassung an beobachtete Daten mitunter nicht so gut möglich ist. In [Wol88, S. 73] wird außerdem kritisiert,

dass aufgrund der Tatsache, dass $E(N) = V(N) = \lambda$ gilt nur wenige auftretende Schadenzahlen wirklich poissonverteilt sein können, wodurch die Poissonverteilung nur eine Näherung darstellen kann.

Allerdings bietet die Poissonverteilung Vorteile bei der Berechnung. Nach [5, S. 20] ist sie gut geeignet, wenn die Wahrscheinlichkeit des Auftretens eines Schadens klein ist und sie kann die Schadenzahl für große Bestände vernünftig beschreiben.

Negative Binomialverteilung

Die negative Binomialverteilung (auch als Pascal-Verteilung bezeichnet) ist eine zweiparametrische Verteilung. Im Gegensatz zur Binomialverteilung können in versicherungsmathematischen Anwendungen beide Parameter zur Modellierung der Schadenzahl verwendet werden.

Definition 3.14.

Eine diskrete Zufallsvariable N heißt negativ-binomialverteilt mit den Parametern r und p (NB(r, p)-verteilt), wenn

$$P(N = k) = \binom{r+k-1}{k} p^r (1-p)^k \quad k = 0, 1, 2, \dots$$

Für die Parameter r und p muss $r \in \mathbb{N}$ und $p \in (0, 1)$ gelten.

Bemerkung. In der Literatur wird die negative Binomialverteilung nicht einheitlich definiert, deswegen findet man auch unterschiedliche Ergebnisse für den Erwartungswert und die Varianz.

Der Erwartungswert und die Varianz der negativen Binomialverteilung lauten:

$$E(N) = \frac{r(1-p)}{p},$$

$$V(N) = \frac{r(1-p)}{p^2}.$$

Da beide Parameter zur Modellierung verwendet werden können, kann die negative Binomialverteilung besser an beobachtete Daten angepasst werden als die Binomialverteilung oder die Poissonverteilung. In [Str88b, S. 116] wird außerdem darauf verwiesen, dass die Gültigkeit von $E(N) < V(N)$ für die negative Binomialverteilung Beobachtungen aus der Praxis in vielen Fällen besser beschreibt als die restriktive Einschränkung der Poissonverteilung, dass $E(N) = V(N)$. Die negative Binomialverteilung eignet sich

dann als Verteilung für die Schadenzahl, wenn die erwartete Anzahl an Schäden kleiner ist als ihre Varianz.

3.2.3 Die Verteilung der Schadenshöhen

Als Verteilungen der Einzelschadenshöhen können folgende, bereits im Kapitel über das individuelle Modell eingeführte Verteilungen verwendet werden:

- Exponentialverteilung,
- Gammaverteilung,
- Inverse Normalverteilung,
- Paretoverteilung,
- Logarithmische Normalverteilung.

Allerdings muss beachtet werden: Eine Verteilung, die sich im individuellen Modell zur Beschreibung eignet muss sich nicht automatisch auch im kollektiven Modell eignen und umgekehrt, denn im individuellen Modell gibt die Verteilung eines Risikos den Jahresgesamtschaden an (der sich aus mehreren Einzelschäden zusammensetzen kann), im kollektiven Modell hingegen werden einzelne Schadenshöhen erfasst.

3.2.4 Berechnung der Verteilung des Gesamtschadens

Die Berechnung der Verteilung des Gesamtschadens ist ein schwieriges, aber für die Praxis bedeutendes Problem. Eine Schwierigkeit liegt darin, dass die Anzahl der Schäden, die in einer Periode eintreten werden, nicht von vornherein klar ist. Eine weitere Schwierigkeit ist die Berechnung der Faltung, die nur für einige Wahrscheinlichkeitsverteilungen explizit möglich ist.

Am Beginn des Kapitels wurde folgende Formel für die Verteilung des Gesamtschadens hergeleitet:

$$F_S(x) = \sum_{n=0}^{\infty} F^{*n}(x) \cdot p_n.$$

Zum expliziten Ausrechnen der Verteilung des Gesamtschadens kann diese Formel aber in den meisten in der Praxis auftretenden Fällen nicht benutzt werden, denn selbst wenn p_n und $F^{*n}(x)$ für alle $n \in \mathbb{N}$ bekannt bzw. überhaupt berechenbar wären müsste der

Grenzwert der Partialsummen von $F^{*n}(x) \cdot p_n$ gebildet werden, was in den meisten Fällen nicht explizit möglich ist (vgl. [Mik09, S. 121]).

Es ist aber notwendig, die Verteilung des Gesamtschadens zu kennen und diese aus den Verteilungen der Schadenzahl und der Schadenshöhen zu berechnen, denn um nur aus beobachteten Daten ein Verteilungsmodell zu erstellen kennt man üblicherweise zu wenige Daten, da der Gesamtschaden nur einmal pro Jahr erfasst wird und sich die Bestände möglicherweise jährlich ändern (vgl. [Mac02, S. 110]).

Zur Ermittlung der Verteilung des Gesamtschadens gibt es folgende Möglichkeiten:

- explizite Berechnung der Verteilung,
- Berechnung mittels numerischer Verfahren,
- Approximation der Verteilung.

Abhängig von den Voraussetzungen ist der Einsatz einer bestimmten Methode nur in gewissen Fällen möglich bzw. sinnvoll. Im Folgenden wird auf alle drei angeführten Methoden eingegangen und es wird beschrieben, wann diese eingesetzt werden können: Zur expliziten Berechnung der Verteilung werden einige Beispiele angeführt, in denen eine explizite Berechnung möglich ist. Zu den numerischen Verfahren wird das Panjer-Verfahren behandelt, das eine rekursive Berechnung erlaubt, wenn die Schadenzahl poisson-, negativ-binomial- oder binomialverteilt ist. Bei der Panjer-Rekursion wird sowohl auf diskrete als auch auf kontinuierliche Schadenshöhenverteilungen eingegangen und es wird gezeigt, wie kontinuierliche Verteilungen diskretisiert werden können. Abschließend wird erörtert, wie das Monte-Carlo-Verfahren zur Approximation von Gesamtschadenverteilungen genutzt werden kann.

Explizite Berechnung der Verteilung des Gesamtschadens

Wie schon erwähnt ist die Berechnung der Gesamtschadenverteilung nur in wenigen Fällen explizit möglich. Ein Fall, in dem diese Verteilung explizit berechnet und auf eine bekannte Verteilung zurückgeführt werden kann, ist folgender (das Beispiel wurde entnommen aus: [Hei87, S.67]):

Sei die typische Schadenshöhe X exponentialverteilt mit dem Parameter λ , dann gilt für die n -fache Faltung, dass sie $\Gamma(n, \lambda)$ -verteilt ist. Das heißt:

$$F^{*n}(x) = \begin{cases} \frac{\Gamma(\lambda x, n)}{(n-1)!} & x > 0, \\ 0 & x \leq 0. \end{cases}$$

Weiters sei N so verteilt dass es sich in der Form $N = M + 1$ darstellen lässt, wobei M negativ-binomialverteilt mit den Parametern $r = 1$ und p (d.h. NB(1, p)-verteilt) ist. Es gilt dann:

$$P(N = k + 1) = \binom{r + k - 1}{k} p^r (1 - p)^k = p(1 - p)^k \quad \text{für } k = 0, 1, 2, \dots;$$

$$P(N = 0) = 0.$$

Dies lässt sich nun in die Formel zur Berechnung der Gesamtschadenverteilung einsetzen. Für $x > 0$ gilt:

$$\begin{aligned} F_S(x) &= \sum_{n=0}^{\infty} F^{*n}(x) \cdot P(N = n) = \\ &= \underbrace{P(N = 0)}_{=0} + \sum_{n=1}^{\infty} \frac{\Gamma(\lambda x, n)}{(n-1)!} p(1-p)^{n-1} = \\ &= \sum_{n=1}^{\infty} p(1-p)^{n-1} \frac{1}{(n-1)!} \int_0^{\lambda x} t^{n-1} e^{-t} dt = \\ &= p \int_0^{\lambda x} e^{-t} \underbrace{\sum_{n=1}^{\infty} \frac{(1-p)^{n-1} t^{n-1}}{(n-1)!}}_{=e^{(1-p)t}} dt = \\ &= p \int_0^{\lambda x} e^{-pt} dt = \\ &= p \left(-\frac{e^{-p\lambda x}}{p} + \frac{1}{p} \right) = \\ &= 1 - e^{-p\lambda x}. \end{aligned}$$

Für $x \leq 0$ gilt:

$$F_S = 0.$$

Der Gesamtschaden ist in diesem Fall also Exp($p\lambda$)-verteilt.

In folgenden Fällen ist es ebenfalls möglich, die Verteilung des Gesamtschadens explizit zu berechnen (entnommen aus: [Sch95, S. 92]):

- Ist N poissonverteilt mit dem Parameter λ und ist X binomialverteilt mit den Parametern $n = 1$ und p , dann ist der Gesamtschaden poissonverteilt mit dem Parameter λp .

- Ist N negativ-binomialverteilt mit den Parametern r und q und ist X binomialverteilt mit den Parametern $n = 1$ und p , dann ist der Gesamtschaden negativ-binomialverteilt mit den Parametern r und $\frac{q}{q+p(1-q)}$.
- Ist N binomialverteilt mit den Parametern m und q und ist X binomialverteilt mit den Parametern $n = 1$ und p , dann ist der Gesamtschaden binomialverteilt mit den Parametern m und $p \cdot q$.

Panjer-Rekursion für diskrete Schadenverteilungen

Ein beliebig genaues Verfahren zur Berechnung der Verteilung des Gesamtschadens stellt das Panjer-Verfahren dar. Es wurde erstmals 1980 vom amerikanischen Mathematiker Harry Panjer publiziert und beruht auf der Idee, dass es Verteilungen der Schadenzahl N gibt, die rekursiv berechnet werden können (vgl. [PW92, S. 194]).

In folgendem Lemma wird gezeigt, dass sich die Wahrscheinlichkeiten poisson-, negativ-binomial- und binomialverteilter Zufallsvariablen rekursiv berechnen lassen:

Lemma 3.4. (i) Sei die Zufallsvariable N poissonverteilt oder negativ-binomialverteilt, dann gibt es $a, b \in \mathbb{R}$ so, dass für alle $n \in \mathbb{N}$ gilt:

$$p_{n+1} = \left(a + \frac{b}{n+1} \right) p_n.$$

(ii) Sei die Zufallsvariable N binomialverteilt mit den Parametern l und p , dann gibt es $a, b \in \mathbb{R}$ so, dass für alle $n \in \mathbb{N}$ mit $n < l$ gilt:

$$p_{n+1} = \left(a + \frac{b}{n+1} \right) p_n.$$

Beweis.

(i) Sei N poissonverteilt mit dem Parameter λ und seien $a := 0$ und $b := \lambda$: Für $n = 0$ gilt: $p_0 = e^{-\lambda}$. Daraus ergibt sich:

$$\left(a + \frac{b}{0+1} \right) p_0 = (0 + \lambda)e^{-\lambda} = \lambda e^{-\lambda} = p_1.$$

Angenommen die Behauptung gilt für ein $n \in \mathbb{N}$, dann gilt für $n+1$:

$$\left(a + \frac{b}{n+1} \right) p_n = \left(0 + \frac{\lambda}{n+1} \right) \frac{\lambda^n}{n!} e^{-\lambda} = \frac{\lambda^{n+1}}{(n+1)!} e^{-\lambda} = p_{n+1}.$$

Sei N negativ-binomialverteilt mit den Parametern r und p und seien $a := 1 - p$ und $b := (r - 1)(1 - p)$: Für $n = 0$ gilt: $p_0 = p^r$. Daraus ergibt sich:

$$\begin{aligned} \left(a + \frac{b}{0 + 1}\right) p_0 &= ((1 - p) + (r - 1)(1 - p))p^r = \\ &= (1 - p)rp^r = \binom{r}{1} p^r (1 - p) = p_1. \end{aligned}$$

Angenommen die Behauptung gilt für ein $n \in \mathbb{N}$, dann gilt für $n + 1$:

$$\begin{aligned} \left(a + \frac{b}{n + 1}\right) p_n &= \left((1 - p) + \frac{(r - 1)(1 - p)}{n + 1}\right) \binom{r + n - 1}{n} p^r (1 - p)^n = \\ &= \frac{(1 - p)(n + 1) + (r - 1)(1 - p)}{n + 1} \binom{r + n - 1}{n} p^r (1 - p)^n = \\ &= \frac{n + r}{n + 1} \binom{r + n - 1}{n} p^r (1 - p)^{n+1} = \\ &= \binom{r + n}{n + 1} p^r (1 - p)^{n+1} = p_{n+1}. \end{aligned}$$

(ii) Sei N binomialverteilt mit den Parametern l und p und seien $a := -\frac{p}{1-p}$ und $b := \frac{(l+1)p}{1-p}$: Für $n = 0$ gilt: $p_0 = (1 - p)^l$. Daraus ergibt sich:

$$\begin{aligned} \left(a + \frac{b}{0 + 1}\right) p_0 &= \left(-\frac{p}{1 - p} + \frac{(l + 1)p}{1 - p}\right) (1 - p)^l = \\ &= \frac{lp}{1 - p} (1 - p)^l = lp(1 - p)^{l-1} = \\ &= \binom{l}{1} p(1 - p)^{l-1} = p_1. \end{aligned}$$

Angenommen die Behauptung gilt für ein $n \in \mathbb{N}$, dann gilt für $n + 1$:

$$\begin{aligned} \left(a + \frac{b}{n + 1}\right) p_n &= \left(-\frac{p}{1 - p} + \frac{(l + 1)p}{(1 - p)(n + 1)}\right) \binom{l}{n} p^n (1 - p)^{l-n} = \\ &= \frac{-p(n + 1) + (l + 1)p}{(1 - p)(n + 1)} \binom{l}{n} p^n (1 - p)^{l-n} = \\ &= \frac{p(l - n)}{(1 - p)(n + 1)} \binom{l}{n} p^n (1 - p)^{l-n} = \end{aligned}$$

$$\begin{aligned}
&= \frac{l-n}{n+1} \binom{l}{n} p^{n+1} (1-p)^{l-n-1} = \\
&= \binom{l}{n+1} p^{n+1} (1-p)^{l-(n+1)} = p_{n+1}.
\end{aligned}$$

□

Satz 3.12. (Panjer-Rekursion für diskrete Schadenverteilungen, vgl. [Mik09, S. 122])
Sei im kollektiven Modell die Schadenzahl N so verteilt, dass es $a, b \in \mathbb{R}$ gibt, sodass

$$p_{n+1} = \left(a + \frac{b}{n+1} \right) p_n \quad \forall n \in \mathbb{N}.$$

Weiters gelte für die identisch verteilten, unabhängigen Schadenshöhen $X_1, X_2, \dots$, dass sie nur Werte in $\mathbb{N}$ annehmen können. Dann gilt für $q_n := P(S = n)$:

$$q_n = \frac{1}{1 - aP(X=0)} \sum_{i=1}^n \left(a + \frac{bi}{n} \right) P(X=i) q_{n-i} \quad \text{für } n \geq 1,$$

$$q_0 = \begin{cases} p_0 & \text{falls } P(X=0) = 0, \\ E(P(X=0)^N) & \text{sonst.} \end{cases}$$

Beweis. nach [Mik09, S. 122ff]

Wie auch schon vorher wird für diesen Beweis die typische Schadenshöhe X verwendet, die die gleiche Verteilung wie die Schadenshöhen $X_1, X_2, \dots$ haben soll, sie wird deshalb als $X := X_n$ für ein passendes n definiert.

Für q_0 gilt:

$$q_0 = \underbrace{P(N=0)}_{=p_0} + P(N > 0, S=0).$$

Ist $P(X=0) = 0$, dann ist $P(N > 0, S=0) = 0$, deshalb gilt in diesem Fall: $q_0 = p_0$.

Ist $P(X=0) \neq 0$, dann gilt:

$$\begin{aligned}
q_0 &= p_0 + \sum_{i=1}^{\infty} P(X_1 = X_2 = \dots = X_i = 0) P(N=i) = \\
&= p_0 + \sum_{i=1}^{\infty} (P(X=0))^i P(N=i) = \\
&= \sum_{i=0}^{\infty} (P(X=0))^i P(N=i) =
\end{aligned}$$

$$= E(P(X=0)^N).$$

Sei $n \geq 1$: Für die Berechnung von q_n gilt dann mit $S_i := \sum_{j=1}^i X_j$:

$$q_n = \sum_{i=1}^{\infty} P(S_i = n) p_i.$$

Nach Voraussetzung gilt für p_i mit $i \geq 1$:

$$p_i = \left(a + \frac{b}{i}\right) p_{i-1}.$$

Daraus ergibt sich für q_n :

$$q_n = \sum_{i=1}^{\infty} P(S_i = n) \left(a + \frac{b}{i}\right) p_{i-1}.$$

Weiters gilt:

$$\begin{aligned} E\left(a + \frac{bX}{n} \middle| S_i = n\right) &= E\left(a + \frac{bX}{S_i} \middle| S_i = n\right) = \\ &= \underbrace{E(a|S_i = n)}_{=a} + E\left(\frac{bX}{S_i} \middle| S_i = n\right) = \\ &= a + bE\left(\frac{X}{S_i} \middle| S_i = n\right) = \\ &= a + \frac{b}{i}E\left(\frac{iX}{S_i} \middle| S_i = n\right) = \\ &= a + \frac{b}{i}E\left(\frac{\sum_{k=1}^i X_k}{S_i} \middle| S_i = n\right) = \\ &= a + \frac{b}{i} \underbrace{E\left(\frac{S_i}{S_i} \middle| S_i = n\right)}_{=1} = a + \frac{b}{i}. \end{aligned}$$

Dieser bedingte Erwartungswert lässt sich auch folgendermaßen berechnen:

$$\begin{aligned}
E\left(a + \frac{bX}{n} \middle| S_i = n\right) &= \sum_{k=0}^n \left(a + \frac{bk}{n}\right) P(X = k | S_i = n) = \\
&= \sum_{k=0}^n \left(a + \frac{bk}{n}\right) \frac{P(X = k)P(S_i = n)}{P(S_i = n)} = \\
&= \sum_{k=0}^n \left(a + \frac{bk}{n}\right) \frac{P(X_i = k)P(S_{i-1} + X_i = n)}{P(S_i = n)} = \\
&= \sum_{k=0}^n \left(a + \frac{bk}{n}\right) \frac{P(X_i = k)P(S_{i-1} = n - k)}{P(S_i = n)} = \\
&= \sum_{k=0}^n \left(a + \frac{bk}{n}\right) \frac{P(X = k)P(S_{i-1} = n - k)}{P(S_i = n)}.
\end{aligned}$$

Durch Gleichsetzen der beiden Ergebnisse für $E(a + \frac{bX}{n} | S_i = n)$ erhält man:

$$a + \frac{b}{i} = \sum_{k=0}^n \left(a + \frac{bk}{n}\right) \frac{P(X = k)P(S_{i-1} = n - k)}{P(S_i = n)}.$$

Dies kann man nun in die Formel für q_n einsetzen:

$$\begin{aligned}
q_n &= \sum_{i=1}^{\infty} P(S_i = n) \left(a + \frac{b}{i}\right) p_{i-1} = \\
&= \sum_{i=1}^{\infty} P(S_i = n) \sum_{k=0}^n \left(a + \frac{bk}{n}\right) \frac{P(X = k)P(S_{i-1} = n - k)}{P(S_i = n)} p_{i-1} = \\
&= \sum_{i=1}^{\infty} \sum_{k=0}^n \left(a + \frac{bk}{n}\right) P(X = k)P(S_{i-1} = n - k) p_{i-1} = \\
&= \sum_{k=0}^n \left(a + \frac{bk}{n}\right) P(X = k) \underbrace{\sum_{i=1}^{\infty} P(S_{i-1} = n - k) p_{i-1}}_{=P(S=n-k)=q_{n-k}} = \\
&= \sum_{k=0}^n \left(a + \frac{bk}{n}\right) P(X = k) q_{n-k} = \\
&= aP(X = 0)q_n + \sum_{k=1}^n \left(a + \frac{bk}{n}\right) P(X = k) q_{n-k}.
\end{aligned}$$

Durch Umformen ergibt sich daraus:

$$q_n - aP(X=0)q_n = \sum_{k=1}^n \left(a + \frac{bk}{n}\right) P(X=k)q_{n-k}$$

$$q_n = \frac{1}{1 - aP(X=0)} \sum_{k=1}^n \left(a + \frac{bk}{n}\right) P(X=k)q_{n-k}.$$

□

Der Satz verlangt, dass die Schadenshöhen in $\mathbb{N}$ liegen, allerdings werden diese üblicherweise mit kontinuierlichen Verteilungen modelliert. Deswegen muss die Verteilung der Schadenshöhen diskretisiert werden. Eine Möglichkeit, die Verteilung zu diskretisieren ist es, nur Werte auf einem Gitter mit der Gitterbreite $h \in \mathbb{R}, h > 0$ zu betrachten. Es ergibt sich dann für die Rekursionsformel:

$$P(S = nh) = \frac{1}{1 - aP(X=0)} \sum_{i=1}^n \left(a + \frac{bi}{n}\right) P(X = ih)P((n-i)h) \quad \text{für } n \geq 1.$$

Je kleiner die Gitterbreite ist, desto genauer kann die Verteilung des Gesamtschadens ermittelt werden, aber es erhöht sich auch der Aufwand zur Berechnung.

Ein Problem bei dieser Art der Diskretisierung ist es, dass der Erwartungswert und die Varianz der Verteilung nicht erhalten bleiben. Bei der Methode des sogenannten local moment matching tritt dieses Problem nicht auf, da die Wahrscheinlichkeitsgewichte nicht auf Gitterpunkte gelegt werden, sondern so, dass der Erwartungswert und die Varianz erhalten bleiben. Das Vorgehen beim local moment matching ist folgendermaßen (nach [10, S. 76]):

Die kontinuierlich verteilte Zufallsvariable X soll durch eine diskrete Verteilung approximiert werden. Sei $h \in \mathbb{R}, h > 0$ die gewählte Schrittweite, dann werden die Wahrscheinlichkeitsgewichte a_i, b_i, c_i mit $i = 0, 2, 4, \dots, K-2$; K eine gerade Zahl, so gewählt, dass:

$$a_i + b_i + c_i = \int_{ih}^{(i+2)h} dF(x),$$

$$(ia_i + (i+1)b_i + (i+2)c_i)h = \int_{ih}^{(i+2)h} x dF(x),$$

$$(i^2a_i + (i+1)^2b_i + (i+2)^2c_i)h^2 = \int_{ih}^{(i+2)h} x^2 dF(x).$$

Man erhält so ein lineares Gleichungssystem in den drei Variablen a_i, b_i und c_i , das immer eindeutig lösbar ist. Weiters soll gelten:

$$\begin{aligned} f_0 &:= a_0, \\ f_k &:= \begin{cases} b_{k-1} & \text{falls } k \text{ ungerade,} \\ c_{k-2} + a_k & \text{falls } k \text{ gerade und } 0 < k < K \end{cases} \\ f_K &:= c_{K-2}. \end{aligned}$$

Es gilt dann:

$$\begin{aligned} \sum_{k=0}^K f_k &= 1, \\ E(X) &= h \sum_{k=0}^K k f_k, \\ E(X^2) &= h^2 \sum_{k=0}^K k^2 f_k. \end{aligned}$$

Damit die dadurch gewonnene diskrete Verteilung die ursprüngliche kontinuierliche Verteilung genau genug approximiert muss K möglichst groß gewählt werden, in versicherungsmathematischen Anwendungen heißt das, dass K größer sein muss als der größtmögliche Schaden, der eintreten kann.

Panjer-Rekursion für kontinuierliche Schadenverteilungen

Im vorhergehenden Abschnitt wurde gezeigt, wie die Verteilung des Gesamtschadens für diskret verteilte Schadenshöhen berechnet werden kann. Um von einer kontinuierlichen Verteilung der Schadenshöhen zu einer diskreten Verteilung zu gelangen musste die Verteilung diskretisiert werden. Dies ist aber nicht immer sinnvoll, deswegen soll in diesem Abschnitt eine Formel zur rekursiven Berechnung der Verteilung des Gesamtschadens behandelt werden, bei der die Diskretisierung nicht notwendig ist. Allerdings führt diese Rekursion zu einer Vollterraschen Integralgleichung, weshalb konkrete Berechnungen verglichen mit der diskreten Variante schwieriger sind.

Im Beweis des nachfolgenden Satzes wird folgende Formel verwendet:

$$\int_0^x \frac{y f(y) f^{*(n-1)}(x-y)}{f^{*n}(x)} dy = \frac{x}{n}.$$

Diese Formel lässt sich folgendermaßen begründen: Die Zufallsvariable X soll die Höhe eines Schadens angeben. Unter der Bedingung, dass n Schäden aufgetreten sind (also $N = n$) und der Gesamtschaden S den Wert x annimmt gilt für den Erwartungswert von X :

$$E(X|S = x, N = n) = \frac{x}{n}.$$

Mit der Formel für bedingte Erwartungswerte lässt sich dieser Erwartungswert auch folgendermaßen berechnen (f bezeichnet dabei die Dichte von X):

$$E(X|S = x, N = n) = \int_0^x \frac{yf(y)f^{*(n-1)}(x-y)}{f^{*n}(x)} dy.$$

Durch Gleichsetzen der beiden Ergebnisse ergibt sich die obige Formel.

Satz 3.13. (*Panjer-Rekursion für kontinuierliche Schadenverteilungen, vgl. [Pan81]*)
 Sei im kollektiven Modell die Schadenzahl N so verteilt, dass es $a, b \in \mathbb{R}$ gibt, sodass

$$p_{n+1} = \left(a + \frac{b}{n+1}\right) p_n \quad \forall n \in \mathbb{N}.$$

Weiters seien die Schadenshöhen $X_1, X_2, \dots$ unabhängig und identisch verteilt mit der Dichte $f(x)$. Dann gilt für die Dichte f_S des Gesamtschadens S :

$$\begin{aligned} f_S(x) &= p_1 f(x) + \int_0^x \left(a + \frac{by}{x}\right) f(y) f_S(x-y) dy \quad \text{für } x > 0, \\ f_S(0) &= p_0. \end{aligned}$$

Beweis. (nach [Pan81, S. 24f])

Für die Dichte des Gesamtschadens gilt:

$$f_S(x) = \sum_{n=0}^{\infty} f^{*n}(x) p_n.$$

Sei $x = 0$, dann gilt:

$$\begin{aligned} f_S(0) &= \sum_{n=0}^{\infty} f^{*n}(0) p_n = \\ &= \underbrace{f^{*0}(0)}_{=1} p_0 + \sum_{n=1}^{\infty} \underbrace{f^{*0}(x)}_{=0} \underbrace{p_n}_{\leq 1} = p_0. \end{aligned}$$

Sei $x > 0$: Es gilt für $f_S(x - y)$ mit $x \neq y$:

$$\begin{aligned} f_S(x - y) &= \sum_{n=0}^{\infty} f^{*n}(x - y)p_n = \\ &= \underbrace{f^{*0}(x - y)}_{=0} p_0 + \sum_{n=1}^{\infty} f^{*n}(x - y)p_n = \\ &= \sum_{n=1}^{\infty} f^{*n}(x - y)p_n. \end{aligned}$$

Das kann in folgenden Ausdruck eingesetzt werden:

$$\begin{aligned} \int_0^x \left(a + \frac{by}{x}\right) f(y) f_S(x - y) dy &= \int_0^x \left(a + \frac{by}{x}\right) f(y) \sum_{n=1}^{\infty} f^{*n}(x - y)p_n dy = \\ &= \sum_{n=1}^{\infty} p_n \int_0^x \left(a + \frac{by}{x}\right) f(y) f^{*n}(x - y) dy = \\ &= \sum_{n=1}^{\infty} p_n \left(\underbrace{a \int_0^x f(y) f^{*n}(x - y) dy}_{=f^{*(n+1)}(x)} + \right. \\ &\quad \left. + \frac{b}{x} \underbrace{\int_0^x y f(y) f^{*n}(x - y) dy}_{=\frac{x}{n+1} f^{*(n+1)}(x)} \right) = \\ &= \sum_{n=1}^{\infty} p_n \underbrace{\left(a + \frac{b}{n+1}\right)}_{=p_{n+1}} f^{*(n+1)}(x) = \\ &= \sum_{n=1}^{\infty} p_{n+1} f^{*(n+1)}(x) = \\ &= \sum_{n=2}^{\infty} p_n f^{*n}(x). \end{aligned}$$

Daraus ergibt sich:

$$\begin{aligned} p_1 f(x) + \int_0^x \left(a + \frac{by}{x}\right) f(y) f_S(x - y) dy &= p_1 f(x) + \sum_{n=2}^{\infty} p_n f^{*n}(x) = \\ &= \sum_{n=1}^{\infty} p_n f^{*n}(x) = \end{aligned}$$

$$= f_S(x).$$

□

Monte-Carlo-Verfahren

Das Monte-Carlo-Verfahren ist eine Vorgangsweise, die auf der Erzeugung von Zufallsgrößen basiert. Durch den Einsatz immer leistungsfähigerer Computer kann in einer kurzen Zeit eine große Zahl an Zufallszahlen generiert werden, was den Einsatz dieser Methode auch dann ermöglicht, wenn eine explizite Ermittlung oder eine numerische Berechnung der Gesamtschadenverteilung nicht möglich ist.

Zur Ermittlung der Verteilung des Gesamtschadens geht man bei bekannter Verteilung der Schadenzahl N und der Schadenshöhen $X_1, X_2, \dots$ folgendermaßen vor (nach [BF87, S. 54] und [Mik09, S. 130]):

1. Entsprechend der Verteilung von N wird eine Folge von zufälligen Größen $N_1, N_2, \dots, N_n$ erzeugt.
2. Zu jedem N_i ($i = 1, 2, \dots, n$) wird entsprechend der Verteilung der Schadenshöhen eine Folge von zufälligen Größen $X_1^{(i)}, X_2^{(i)}, \dots, X_{N_i}^{(i)}$ generiert.
3. Die Summe $S^{(i)} = \sum_{j=1}^{N_i} X_j^{(i)}$ wird für jedes $i = 1, 2, \dots, n$ berechnet, dadurch erhält man die Folge $S^{(1)}, S^{(2)}, \dots, S^{(n)}$.
4. Sei $m(x)$ die Anzahl aller $S^{(i)}$ mit $S^{(i)} \leq x$, dann gilt für die Verteilungsfunktion des Gesamtschadens:

$$F_S(x) \approx \frac{m(x)}{n}.$$

Als Vorteile dieser Methode werden in [BF87, S. 55] genannt, dass das Verfahren einfach und flexibel ist, dass auch empirisch ermittelte Verteilungen im Verfahren berücksichtigt werden können und dass Auswirkungen getroffener Maßnahmen auf einfache Weise simuliert und beobachtet werden können. Demgegenüber werden die Nachteile gestellt, dass Computer keine „echten“ Zufallszahlen generieren können, sondern nur Pseudozufallszahlen, die nach einem gewissen Algorithmus berechnet werden. Außerdem kann das Verfahren nur Ergebnisse liefern, die mit einer gewissen Wahrscheinlichkeit richtig sind, je höher diese Wahrscheinlichkeit sein soll, desto mehr Zahlen müssen generiert werden. Weiters kann es bei Beständen, in denen Großschäden auftreten können, zu Problemen kommen, da diese Großschäden mitunter nicht ausreichend berücksichtigt werden.

In [Mik09, S. 131, S. 136] wird darauf verwiesen, dass das Monte-Carlo-Verfahren und aus der Idee des Monte-Carlo-Verfahrens entwickelte Verfahren durch die enorme Leistungssteigerung bei Computern in den letzten Jahrzehnten stark an Bedeutung gewonnen haben. Dennoch wird dargelegt, dass diese Methode nicht perfekt ist und dass sie deshalb vor allem dann eingesetzt werden soll, wenn kein anderer Weg zum Ziel führt. Das Monte-Carlo-Verfahren eignet sich vor allem für kleine Schadeneintrittswahrscheinlichkeiten und mittlere Bestandsgrößen.

4 Die Berechnung von Prämien

Nachdem beschrieben wurde, wie Versicherungen Modelle mit den ihnen anvertrauten Risiken erstellen soll nun die Prämienberechnung behandelt werden. Bei der Berechnung von Versicherungsprämien wird das Ziel verfolgt, einen fairen Betrag zu ermitteln, den VersicherungsnehmerInnen dem Versicherungsunternehmen zu zahlen haben. Was „fair“ in diesem Zusammenhang bedeuten soll kann aber nicht eindeutig bestimmt werden. So wäre für VersicherungsnehmerInnen das Äquivalenzprinzip fair, nach dem die Prämie genau so hoch ist wie der zu erwartende Schaden. Jedoch gilt dann, wie später gezeigt wird, dass das Versicherungsunternehmen mit einer Wahrscheinlichkeit von 50% nach nur einer betrachteten Periode mehr Geld auszahlen musste, als es eingenommen hat. Außerdem entstehen für das Versicherungsunternehmen Kosten, die gedeckt werden müssen. So kann das Äquivalenzprinzip, konsequent zu Ende gedacht, für VersicherungsnehmerInnen nicht zielführend sein, da das Versicherungsunternehmen in diesem Fall mit hoher Wahrscheinlichkeit irgendwann nicht mehr zahlungsfähig sein wird.

Im Gegensatz zu VersicherungsnehmerInnen hat das Versicherungsunternehmen das Interesse, möglichst hohe Prämien anzusetzen, damit alle Schäden und die sonstigen Kosten beglichen werden können und mit hoher Wahrscheinlichkeit ein Gewinn verbucht werden kann. Allerdings werden sich kaum VersicherungsnehmerInnen finden, die bereit sind, überhöhte Prämien zu zahlen, besonders dann nicht, wenn es am Markt weitere Unternehmen gibt, die Versicherungsleistungen zu niedrigeren Preisen anbieten.

In diesem Kapitel werden zwei voneinander verschiedene Methoden zur Prämienbestimmung behandelt. Die erste Methode ist die Prämienberechnung mittels Prämienprinzipien: Hierbei werden wahrscheinlichkeitstheoretische Eigenschaften von Risiken (oder ganzen Beständen), wie beispielsweise der Erwartungswert und die Varianz von Schadenshöhen zur Prämienermittlung herangezogen. Bei der Erfahrungstarifizierung hingegen betrachtet man den individuellen Schadenverlauf eines Risikos: Die Prämienhöhe wird bei dieser Art der Prämienberechnung davon abhängig gemacht, wie viele Schäden in vorhergehenden Perioden aufgetreten sind bzw. wie hoch diese Schäden waren.

4.1 Grundlegende Eigenschaften von Beständen

Zuerst sollen zwei Eigenschaften genannt werden, die einen Bestand von Risiken auszeichnen und als Grundlage für die Berechnung von Prämien dienen, nämlich die Stabilität in der Größe und die Stabilität in der Zeit (die Beschreibung dieser beiden Eigenschaften erfolgt nach [Kre85, S. 20f]).

Stabilität in der Größe

Sei $\{X_1, X_2, X_3, \dots\}$ ein Bestand von unabhängigen Risiken. Falls der Grenzwert

$$\nu := \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n E(X_i)$$

existiert, dann gilt nach dem starken Gesetz der großen Zahlen:

$$\frac{1}{n} \sum_{i=1}^n X_i \rightarrow \nu \quad \text{fast sicher.}$$

Das heißt, dass in einem großen Bestand für jedes Risiko der Schaden ν erwartet werden kann.

Stabilität in der Zeit

Nun soll der Bestand über k Jahre betrachtet werden. Sei $\{X_{1,j}, X_{2,j}, \dots, X_{n,j}\}$ ein Bestand von unabhängigen Risiken im Jahr j ($j \in \{1, 2, 3, \dots\}$). Falls der Grenzwert

$$\mu_i := \lim_{k \rightarrow \infty} \frac{1}{k} \sum_{j=1}^k E(X_{i,j})$$

für alle $i \in \{1, 2, \dots, n\}$ existiert, dann gilt nach dem starken Gesetz der großen Zahlen für alle $i \in \mathbb{N}$:

$$\frac{1}{k} \sum_{j=1}^k X_{i,j} \rightarrow \mu_i \quad \text{fast sicher.}$$

Das heißt, dass für das Risiko X_i in jedem Jahr der mittlere Schaden μ_i erwartet werden kann.

4.2 Prämienberechnung mit Prämienprinzipien

Unter einem Prämienprinzip versteht man ein Funktional, das jedem Risiko eine Prämie zuordnet, es kann folgendermaßen definiert werden (in der Literatur findet man verschiedene Definitionen eines Prämienprinzips, die hier gegebene Definition folgt der Definition in [Sch02, S. 240]):

Definition 4.1.

Sei $X \subseteq \{Y | Y \text{ ist ein Risiko}\}$. Als Prämienprinzip wird eine Abbildung

$$H : X \rightarrow (0, \infty)$$

bezeichnet, wenn sie folgende Eigenschaften besitzt:

1. Seien die Risiken X, Y identisch verteilt, dann gilt: $H(X) = H(Y)$.
2. Für jedes Risiko X gilt: $E(X) \leq H(X)$.

Nach den geforderten Eigenschaften muss die Prämie also für identisch verteilte Risiken gleich hoch sein und sie muss mindestens so hoch sein wie der Erwartungswert des Risikos. Dass für identisch verteilte Risiken gleich hohe Prämien ermittelt werden sollen ist im Sinne der „Fairness“ klar. Außerdem ist es leicht einsichtig, dass die Prämie nicht niedriger sein kann als der erwartete Schaden.

Bemerkung. In [Sch02] wird zusätzlich gefordert, dass für ein Prämienprinzip H die Bedingung $P(X > H(X)) > 0$ gelten muss (wobei X ein beliebiges Risiko darstellt). Diese Bedingung wird als No-Abitrage-Bedingung bezeichnet und sorgt dafür, dass das Versicherungsunternehmen keinen sicheren Gewinn beim Abschluss eines Versicherungsvertrages erwarten kann, denn wenn die No-Abitrage-Bedingung gilt, dann ist es möglich, dass Schäden eintreten, deren Schadenssumme höher als die Prämie ist (vgl. [Sch02, S. 240]). Hier wird darauf verzichtet, die Gültigkeit von $P(X > H(X)) > 0$ für ein Prämienprinzip zu verlangen, da diese in vielen Fällen von Parametern des Prämienprinzips (wie beispielsweise dem Sicherheitszuschlag) abhängt. Sind die konkreten Parameter eines Prämienprinzips bekannt, dann ist die Überprüfung der No-Abitrage-Bedingung aber sinnvoll.

In dieser Arbeit sollen nur Risiken als versicherbar angesehen werden, die nicht konstant sind (für konstante Risiken wäre der Abschluss eines Versicherungsvertrages auch

nicht sinnvoll). Ist ein Risiko X nicht konstant, dann muss $E(X) > 0$ und $V(X) > 0$ gelten. Außerdem sollen aus naheliegenden Gründen für ein versicherbares Risiko X die Bedingungen $E(X) < \infty$ und $V(X) < \infty$ gelten.

Definition 4.2.

Ein Risiko X heißt versicherbar, wenn $0 < E(X) < \infty$ und $0 < V(X) < \infty$ gilt.

Es sollen nun einige wünschenswerte Eigenschaften von Prämienprinzipien definiert werden:

Definition 4.3.

Sei H ein Prämienprinzip und seien X und Y zwei beliebige, unabhängige Risiken.

1. H heißt additiv, falls $H(X + Y) = H(X) + H(Y)$ für alle versicherbaren Risiken X, Y gilt.
2. Gilt $H(X + a) = H(X) + a$ für alle $a \in \mathbb{R}$ mit $a \geq 0$ dann wird H als translationsinvariant bezeichnet.
3. Falls $H(aX) = aH(X)$ für alle $a \in \mathbb{R}$ mit $a \geq 0$ gilt, dann heißt H positiv homogen.
4. H wird maximalschadenbegrenzt genannt, falls

$$H(X) \leq \inf\{x \in \mathbb{R} : P(X \leq x) = 1\}.$$

Diese Eigenschaften lassen sich folgendermaßen interpretieren (nach [5, S. 33f], [13, S. 54f] und [11, S. 1]):

- Die Additivität bedeutet, dass es weder durch das Zusammenfügen, noch durch das Trennen von (unabhängigen) Risiken Vorteile für VersicherungsnehmerInnen oder für Versicherungsunternehmen gibt.
- Ein translationsinvariantes Prämienprinzip besitzt die Eigenschaft, dass sich bei Erhöhung des Risikos um einen festen Betrag die Prämie um den selben Betrag erhöht.
- Ist ein Prämienprinzip positiv homogen, dann bedeutet das, dass die Prämie direkt proportional zum Risiko ist, betrachtet man also ein a -faches Risiko, dann ergibt sich für dieses eine a -fache Prämie.

- Bei einem maximalschadenbegrenzten Prämienprinzip ist die Versicherungsprämie stets kleiner als der größtmögliche Schaden. Gilt die Maximalschadenbegrenztheit nicht, dann macht es für VersicherungsnehmerInnen keinen Sinn, eine Versicherung abzuschließen. (Wenn die No-Abitrage-Bedingung $P(X > H(X)) > 0$ für ein Prämienprinzip gilt, dann folgt daraus, dass es auch maximalschadenbegrenzt ist.)
- Bei einem positiv homogenen Prämienprinzip und bei einem additiven Prämienprinzip wird außerdem die Möglichkeit eines risikolosen Gewinnes für das Versicherungsunternehmen ausgeschlossen, denn wäre $H(X + Y) > H(X) + H(Y)$ bzw. $H(aX) > aH(X)$, dann würde das Versicherungsunternehmen einen Vorteil daraus ziehen.

4.2.1 Das Nettoprämienprinzip

Als erstes Prämienprinzip wird das Nettoprämienprinzip behandelt. Nach diesem wird die Prämie genau so hoch angesetzt wie die Höhe des zu erwartenden Schadens:

Definition 4.4.

Sei X ein versicherbares Risiko und H ein Prämienprinzip, das durch

$$H(X) := E(X)$$

definiert ist. Dann wird H als Nettoprämienprinzip bezeichnet.

Satz 4.1. *Das Nettoprämienprinzip ist additiv, translationsinvariant, positiv homogen und maximalschadenbegrenzt.*

Beweis.

Seien X und Y zwei unabhängige, versicherbare Risiken und sei $a \in \mathbb{R}$ beliebig mit $a \geq 0$. Dann gilt:

Additivität:

$$H(X + Y) = E(X + Y) = E(X) + E(Y) = H(X) + H(Y).$$

Translationsinvarianz:

$$H(X + a) = E(X + a) = E(X) + a = H(X) + a.$$

Positive Homogenität:

$$H(aX) = E(aX) = aE(X) = aH(X).$$

Maximalschadenbegrenztheit: Sei $u := \inf\{x \in \mathbb{R} : P(X \leq x) = 1\}$. Für jedes versicherbare Risiko X gilt: $X \leq u$. Deshalb gilt: $E(X) \leq E(u) = u$. Das heißt für $H(X)$:

$$H(X) = E(X) \leq \inf\{x \in \mathbb{R} : P(X \leq x) = 1\}.$$

□

Das Nettoprämienprinzip lässt sich durch die Stabilität in der Größe und die Stabilität in der Zeit begründen. Allerdings ist dieses Prämienprinzip für die Praxis nicht relevant, denn für real abgeschlossene Versicherungsverträge gilt (nach [Kre85, S. 21]):

1. Die Bestände eines Versicherungsunternehmens umfassen nur endlich viele Risiken.
2. Die Gültigkeit eines Versicherungsvertrages ist zeitlich begrenzt.
3. Versicherungsunternehmen müssen nicht nur aufgetretene Schäden begleichen, sondern es fallen auch Kosten für Mitarbeiter, Verwaltung etc. an. Außerdem verfolgen Versicherungsunternehmen das Ziel, einen Gewinn zu erzielen.

Weiters kann gezeigt werden, dass das Versicherungsunternehmen mit relativ hoher Wahrscheinlichkeit nicht so viel einnimmt wie es für die Begleichung von Schäden ausgeben muss: Sei $X_1, X_2, \dots, X_n$ ein Bestand von unabhängigen, identisch verteilten Risiken mit $\mu := E(X_i)$ und $\sigma^2 = V(X_i)$ für $i = 1, 2, \dots, n$. Sei weiters $S = \sum_{i=1}^n X_i$ der Gesamtschaden in einer Betrachtungsperiode. Nach dem zentralen Grenzwertsatz gilt dann:

$$\lim_{n \rightarrow \infty} P\left(\frac{\sum_{i=1}^n X_i - n\mu}{\sqrt{n\sigma^2}} \leq x\right) = \Phi(x) \quad \forall x \in \mathbb{R}.$$

Für $x = 0$ gilt:

$$P\left(\frac{\sum_{i=1}^n X_i - n\mu}{\sqrt{n\sigma^2}} \leq 0\right) = \Phi(0) = \frac{1}{2}.$$

Das heißt, mit einer Wahrscheinlichkeit von 50% muss das Versicherungsunternehmen mindestens so viel auszahlen wie es eingenommen hat.

Es ist also offensichtlich, dass die Nettoprämie nicht ausreicht, um über längere Zeit kostendeckend zu arbeiten. Das Nettoprämienprinzip dient aber als Ausgangspunkt für die Definition weiterer Prämienprinzipien.

Bemerkung. In der Lebensversicherung ist das Nettoprämienprinzip das Standardprinzip zur Berechnung von Versicherungsprämien, allerdings werden dort die Daten (beispielsweise Sterbetafeln) so angepasst, dass die berechnete Nettoprämie bereits einen Zuschlag enthält (vgl. [5, S. 31]).

4.2.2 Aus dem Nettoprämienprinzip abgeleitete Prinzipien

Die Bruttoprämie, die VersicherungsnehmerInnen zu zahlen haben, muss gemäß den getätigten Überlegungen höher sein als die Nettoprämie. Sie setzt sich aus folgenden Teilen zusammen (nach [10, S. 13]):

- Nettoprämie $E(X)$ für ein Risiko X ,
- Sicherheitszuschlag (manchmal auch als Schwankungszuschlag bezeichnet),
- Betriebskosten- und Gewinnzuschlag,
- Steuern.

Die Zuschläge von Betriebskosten, Steuern und dem Gewinn werden hier nicht näher betrachtet, da diese Zuschläge üblicherweise nicht von VersicherungsmathematikerInnen kalkuliert werden. Es soll aber auf die Risikoprämie eingegangen werden, die sich aus der Nettoprämie und einem Sicherheitszuschlag $c \in \mathbb{R}$, $c \geq 0$ zusammensetzt.

Im Folgenden werden einige Prämienprinzipien beschrieben, die aus dem Nettoprämienprinzip abgeleitet werden können.

Das Erwartungswertprinzip

Beim Erwartungswertprinzip wird nicht der Erwartungswert eines Risikos als Prämie herangezogen, sondern ein höherer, zum Erwartungswert proportionaler Betrag.

Definition 4.5.

Sei X ein versicherbares Risiko und H ein Prämienprinzip, das durch

$$H(X) := E(X) + cE(X) = (1 + c)E(X) \quad c \in \mathbb{R}, c > 0$$

definiert ist. Dann wird H als Erwartungswertprinzip bezeichnet.

Satz 4.2. *Das Erwartungswertprinzip ist additiv, positiv homogen, nicht translationsinvariant und nicht maximalschadenbegrenzt.*

Beweis.

Seien X und Y zwei unabhängige, versicherbare Risiken und sei $a \in \mathbb{R}$ beliebig mit $a > 0$. Dann gilt für das Erwartungswertprinzip bezüglich der zu untersuchenden Eigenschaften:
Additivität:

$$\begin{aligned} H(X + Y) &= (1 + c)E(X + Y) = (1 + c)(E(X) + E(Y)) = \\ &= (1 + c)E(X) + (1 + c)E(Y) = H(X) + H(Y). \end{aligned}$$

Positive Homogenität:

$$H(aX) = (1 + c)E(aX) = a(1 + c)E(X) = aH(X).$$

Translationsinvarianz:

$$\begin{aligned} H(X + a) &= (1 + c)E(X + a) = (1 + c)E(X) + (1 + c)a = \\ &= H(X) + (1 + c)a \neq H(X) + a. \end{aligned}$$

Maximalschadenbegrenztheit: Sei $b \in \mathbb{R}, b > 0$ und sei $P(X = b) = P(X = 2b) = \frac{1}{2}$, dann gilt:

$$\begin{aligned} E(X) &= \frac{1}{2}b + \frac{1}{2}2b = \frac{3}{2}b \\ 2b &= \inf\{x \in \mathbb{R} : P(X \leq x) = 1\}. \end{aligned}$$

Das heißt:

$$H(X) = (1 + c)E(X) = \frac{3}{2}(1 + c)b.$$

Deshalb gilt:

$$H(X) > \inf\{x \in \mathbb{R} : P(X \leq x) = 1\} \Leftrightarrow \frac{3}{2}(1 + c)b > 2b \Leftrightarrow c > \frac{1}{3}.$$

Somit ist das Erwartungswertprinzip nicht maximalschadenbegrenzt. □

In [13, S. 48] wird darauf eingegangen, dass bei diesem Prämienprinzip nur der Erwartungswert maßgeblich für die Berechnung der Prämie ist, die Varianz hingegen wird nicht berücksichtigt. Dies kann sowohl für das Versicherungsunternehmen, als auch für VersicherungsnehmerInnen nachteilig sein: Ist die Varianz klein, dann erzielt das Versicherungsunternehmen mit großer Wahrscheinlichkeit einen Gewinn, ist die Varianz hin-

gegen groß, dann können hohe Schäden auftreten und das Versicherungsunternehmen muss mit einem Verlust rechnen.

Das Varianzprinzip

Im Gegensatz zum Erwartungswertprinzip wird beim Varianzprinzip auch die Varianz eines Risikos in die Prämienberechnung einbezogen, Schwankungen des Risikos werden also berücksichtigt.

Definition 4.6.

Sei X ein versicherbares Risiko und H ein Prämienprinzip, das durch

$$H(X) := E(X) + cV(X) \quad c \in \mathbb{R}, c > 0$$

definiert ist. Dann wird H als Varianzprinzip bezeichnet.

Satz 4.3. *Das Varianzprinzip ist additiv, translationsinvariant, nicht positiv homogen und nicht maximalschadenbegrenzt.*

Beweis. nach [5, S. 35]

Seien X und Y zwei unabhängige, versicherbare Risiken und sei $a \in \mathbb{R}$ beliebig mit $a > 0$. Dann gilt für das Varianzprinzip bezüglich der zu untersuchenden Eigenschaften: Additivität:

$$\begin{aligned} H(X + Y) &= E(X + Y) + cV(X + Y) \\ &= E(X) + E(Y) + cV(X) + cV(Y) = H(X) + H(Y) \end{aligned}$$

Translationsinvarianz:

$$H(X + a) = E(X + a) + c \underbrace{V(X + a)}_{=V(X)} = E(X) + a + cV(X) = H(X) + a$$

Positive Homogenität:

$$\begin{aligned} H(aX) &= E(aX) + cV(aX) = aE(X) + a^2cV(X) \neq \\ &\neq aE(X) + acV(X) = aH(X) \quad \forall a \neq 1. \end{aligned}$$

Maximalschadenbegrenztheit: Sei $b \in \mathbb{R}, b > 0$ und $P(X = 1+b) = \alpha, P(X = 1) = 1 - \alpha$

für ein $\alpha \in (0, 1)$, dann gilt:

$$\begin{aligned} E(X) &= (1+b)\alpha + (1-\alpha) = 1+b\alpha, \\ V(X) &= E(X^2) - (E(X))^2 = (1+b)^2\alpha + (1-\alpha) - (1+b\alpha)^2 = \\ &= b^2\alpha(1-\alpha). \end{aligned}$$

Das heißt:

$$H(X) = E(X) + cV(X) = 1+b\alpha + cb^2\alpha(1-\alpha).$$

Außerdem gilt:

$$\inf\{x \in \mathbb{R} : P(X \leq x) = 1\} = 1+b.$$

Deshalb gilt:

$$H(X) > \inf\{x \in \mathbb{R} : P(X \leq x) = 1\} \Leftrightarrow 1+b\alpha + cb^2\alpha(1-\alpha) > 1+b.$$

Setzt man den oben ermittelten Wert für $H(X)$ ein dann heißt das:

$$\begin{aligned} 1+b\alpha + cb^2\alpha(1-\alpha) &> 1+b \\ \alpha + cb\alpha(1-\alpha) &> 1 \\ b &> c\alpha. \end{aligned}$$

Wird also $b > c\alpha$ gewählt, dann gilt:

$$H(X) = 1+b\alpha + cb^2\alpha(1-\alpha) > 1+b = \inf\{x \in \mathbb{R} : P(X \leq x) = 1\}.$$

Das Varianzprinzip ist also nicht maximalschadenbegrenzt. □

Durch das Einbeziehen der Varianz in die Prämienberechnung hat das Varianzprinzip Vorteile gegenüber den bisher behandelten Prinzipien, allerdings wird in [13, S. 48] angemerkt, dass dieses Prinzip als nicht fair gegenüber VersicherungsnehmerInnen betrachtet werden muss, da die (positive) Varianz zum Erwartungswert addiert wird. So gehen Abweichungen unterhalb des Erwartungswertes genau so stark ein wie Abweichungen überhalb des Erwartungswertes. Für das Versicherungsunternehmen bedeuten Abweichungen unterhalb des Erwartungswertes aber, dass der Schaden kleiner als der Erwartungswert des Risikos ist und dass deshalb ein geringerer Geldbetrag als der Er-

wartungswert ausbezahlt werden muss.

Außerdem werden bei diesem Prinzip Größen verschiedener Dimension addiert, was unvorteilhaft ist.

Das Standardabweichungsprinzip

Beim Standardabweichungsprinzip wird wie beim Varianzprinzip berücksichtigt, wie stark ein Risiko streut. Allerdings wird hier nicht ein Vielfaches der Varianz, sondern ein Vielfaches der Quadratwurzel der Varianz addiert.

Definition 4.7.

Sei X ein versicherbares Risiko und H ein Prämienprinzip, das durch

$$H(X) := E(X) + c\sqrt{V(X)} \quad c \in \mathbb{R}, c > 0$$

definiert ist. Dann wird H als Standardabweichungsprinzip bezeichnet.

Satz 4.4. *Das Standardabweichungsprinzip ist positiv homogen, translationsinvariant, nicht additiv und nicht maximalschadenbegrenzt.*

Beweis. nach [5, S. 35]

Seien X und Y zwei unabhängige, versicherbare Risiken und sei $a \in \mathbb{R}$ beliebig mit $a > 0$. Dann gilt für das Standardabweichungsprinzip bezüglich der zu untersuchenden Eigenschaften:

Positive Homogenität:

$$\begin{aligned} H(aX) &= E(aX) + c\sqrt{V(aX)} = aE(X) + c\sqrt{a^2V(X)} = \\ &= aE(X) + ac\sqrt{V(X)} = aH(X). \end{aligned}$$

Translationsinvarianz:

$$\begin{aligned} H(X + a) &= E(X + a) + c\underbrace{\sqrt{V(X + a)}}_{=\sqrt{V(X)}} = \\ &= E(X) + a + c\sqrt{V(X)} = H(X) + a. \end{aligned}$$

Additivität:

$$\begin{aligned} H(X+Y) &= E(X+Y) + c\sqrt{V(X+Y)} \neq \\ &\neq E(X) + c\sqrt{V(X)} + E(Y) + c\sqrt{V(Y)} = H(X) + H(Y). \end{aligned}$$

Maximalschadenbegrenztheit: Sei $b \in \mathbb{R}, b > 0$ und $P(X = 1+b) = \alpha, P(X = 1) = 1-\alpha$ für ein $\alpha \in (0, 1)$, dann gilt nach dem Beweis zum Varianzprinzip:

$$\begin{aligned} E(X) &= 1 + b\alpha, \\ V(X) &= b^2\alpha(1-\alpha). \end{aligned}$$

Das heißt:

$$H(X) = E(X) + c\sqrt{V(X)} = 1 + b\alpha + cb\sqrt{\alpha(1-\alpha)}.$$

Außerdem gilt:

$$\inf\{x \in \mathbb{R} : P(X \leq x) = 1\} = 1 + b.$$

Deshalb gilt:

$$H(X) > \inf\{x \in \mathbb{R} : P(X \leq x) = 1\} \Leftrightarrow 1 + b\alpha + cb\sqrt{\alpha(1-\alpha)} > 1 + b.$$

Setzt man den oben ermittelten Wert für $H(X)$ ein dann heißt das:

$$\begin{aligned} 1 + b\alpha + cb\sqrt{\alpha(1-\alpha)} &> 1 + b \\ \alpha + c\sqrt{\alpha(1-\alpha)} &> 1 \\ c &> \frac{1-\alpha}{\sqrt{\alpha(1-\alpha)}} = \frac{\sqrt{1-\alpha}}{\sqrt{\alpha}}. \end{aligned}$$

Ist also $c > \frac{\sqrt{1-\alpha}}{\sqrt{\alpha}}$, dann gilt:

$$H(X) = 1 + b\alpha + cb\sqrt{\alpha(1-\alpha)} > 1 + b = \inf\{x \in \mathbb{R} : P(X \leq x) = 1\}.$$

Das Standardabweichungsprinzip ist also nicht maximalschadenbegrenzt. □

Das Problem, dass nur positive Beträge zum Erwartungswert addiert werden, ist beim Standardabweichungsprinzip wie beim Varianzprinzip vorhanden. Allerdings hat die Wurzel der Varianz hier die gleiche Dimension wie der Erwartungswert.

4.2.3 Berechnung von Prämien mit Nutzenfunktionen

Nun soll versucht werden, den Nutzen, den ein Versicherungsunternehmen durch das Abschließen eines Versicherungsvertrages hat, in die Prämienberechnung miteinzubeziehen. Dieser Nutzen soll durch eine Funktion, die Nutzenfunktion, dargestellt werden. Die Nutzenfunktion gibt an, wie hoch der Nutzen eines gewissen Geldbetrages für das Versicherungsunternehmen ist.

Definition 4.8.

Sei $v : \mathbb{R} \rightarrow \mathbb{R}$ eine zweimal differenzierbare Funktion. Dann wird diese als Nutzenfunktion bezeichnet, falls:

$$\begin{aligned} v'(x) &> 0 & \forall x \in \mathbb{R}, \\ v''(x) &\leq 0 & \forall x \in \mathbb{R}. \end{aligned}$$

Diese beiden Bedingungen für eine Nutzenfunktion kann man folgendermaßen interpretieren (vgl. [13, S. 49]): Da eine Nutzenfunktion streng monoton wachsend sein muss ist der Nutzen höher, wenn das Kapital zunimmt. Wenn das Kapital konstant bleibt steigt der Nutzen nicht. Die Eigenschaft der Konkavität bedeutet, dass eine Zunahme eines niedrigen Kapitals höher bewertet wird als die Zunahme eines hohen Kapitals.

Das allgemeine Nullnutzenprinzip

Definition 4.9.

Sei X ein versicherbares Risiko und v eine Nutzenfunktion. Sei weiters $k \in \mathbb{R}$ die Anfangsrisikoreserve, d. h. jener Betrag, der am Beginn der Versicherungsperiode zur Begleichung von Schäden des Risikos X zur Verfügung steht. Das Prämienprinzip $H(X)$, das die Gleichung

$$v(k) = E(v(k + H(X) - X))$$

erfüllt wird Nullnutzenprinzip genannt.

Dieses Prinzip beruht auf folgender Überlegung (vgl. [4, S. 9]): Der Funktionswert $v(k)$ gibt den Nutzen des Betrages k an. Die Nullnutzenprämie $H(X)$ gibt deswegen jenen Betrag an, bei dem der erwartete Nutzen durch das Versichern des Risikos X gleich groß ist wie der Nutzen der Anfangsrisikoreserve k . Für das Versicherungsunternehmen stellt $H(X)$ den kleinsten Betrag dar, für den sich eine Versicherung des Risikos X lohnt.

Das Nullnutzenprinzip kann in dem Sinn, dass weder für VersicherungsnehmerInnen, noch für das Versicherungsunternehmen ein Vorteil entsteht, als fair betrachtet werden.

Lemma 4.1. *Sei X ein versicherbares Risiko, v eine Nutzenfunktion und $k \in \mathbb{R}$. Besitzt die Gleichung*

$$v(k) = E(v(k + H(X) - X))$$

eine Lösung $H(X)$, dann ist diese Lösung eindeutig.

Beweis. nach [9, S. 3]

Angenommen die Gleichung würde zwei Lösungen h_1, h_2 mit $h_1 < h_2$ besitzen. Da die Nutzenfunktion v streng monoton wachsend ist gilt dann:

$$v(k + h_1 - X) < v(k + h_2 - X).$$

Deshalb gilt auch:

$$E(v(k + h_1 - X)) < E(v(k + h_2 - X)).$$

Da aber $E(v(k + h_1 - X)) = v(k) = E(v(k + h_2 - X))$ gelten muss stellt das einen Widerspruch dar. $\square$

Satz 4.5. *Das Nullnutzenprinzip ist translationsinvariant, maximalschadenbegrenzt, nicht positiv homogen und nicht additiv.*

Beweis. nach [Wol88, 212f]

Seien X und Y zwei unabhängige, versicherbare Risiken. Dann gilt bezüglich der zu untersuchenden Eigenschaften:

Translationsinvarianz: Für die Prämie $H(X)$ gilt:

$$v(k) = E(v(k + H(X) - X)).$$

Außerdem gilt für $H(X + a)$ mit $a > 0$:

$$v(k) = E(v(k + H(X + a) - (X + a))) = E(v(k + (H(X + a) - a) - X)).$$

Da die Lösung der Gleichung $v(k) = E(v(k + H(X) - X))$ nach dem obigen Lemma eindeutig sein muss, muss deshalb gelten:

$$H(X) = H(X + a) - a.$$

Das heißt:

$$H(X + a) = H(X) + a.$$

Maximalschadenbegrenztheit: Sei $x_m := \inf\{x \in \mathbb{R} : P(X \leq x) = 1\}$, dann gilt für das Risiko X : $X \leq x_m$. Deshalb gilt auch:

$$k + H(X) - X \geq k + H(X) - x_m.$$

Da die Nutzenfunktion v streng monoton wachsend ist folgt daraus:

$$E(v(k + H(X) - X)) \geq E(v(k + H(X) - x_m)).$$

Für $E(v(k + H(X) - X))$ muss $E(v(k + H(X) - X)) = v(k)$ gelten, außerdem gilt $E(v(k + H(X) - x_m)) = v(k + H(X) - x_m)$. Deshalb muss folgendes gelten:

$$v(k) \geq v(k + H(X) - x_m).$$

Aufgrund der strengen Monotonie der Nutzenfunktion v ergibt sich daraus:

$$\begin{aligned} k &\geq k + H(X) - x_m, \\ 0 &\geq H(X) - x_m. \end{aligned}$$

Das heißt:

$$H(X) \leq x_m = \inf\{x \in \mathbb{R} : P(X \leq x) = 1\}.$$

Additivität: Dass das Nullnutzenprinzip nicht additiv ist soll durch ein Gegenbeispiel gezeigt werden: Sei das Risiko X folgendermaßen verteilt:

$$P(X = 0) = P(X = 1) = \frac{1}{2}.$$

Die Nutzenfunktion v wird folgendermaßen definiert:

$$v(x) := \log(x + 1).$$

Sei außerdem Y ein Risiko, das identisch verteilt ist wie das Risiko X , dann gilt:

$$P(X + Y = 0) = P(X + Y = 2) = \frac{1}{4},$$

$$P(X + Y = 1) = \frac{1}{2}.$$

Weiters soll $k := 0$ gelten. Dann gilt für $H(X)$:

$$\begin{aligned} v(k) &= E(v(k + H(X) - X)) \\ \log(1) &= E(\log(H(X) - X + 1)) \\ 0 &= \frac{1}{2} \log(H(X) + 1) + \frac{1}{2} \log(H(X)) \\ 0 &= \log(H(X) + 1) + \log(H(X)) \\ 0 &= \log((H(X) + 1)H(X)) \\ 1 &= (H(X) + 1)H(X) \\ 0 &= (H(X))^2 + H(X) - 1. \end{aligned}$$

Da nur positive Werte als Prämie zugelassen sind erfüllt nur $H(X) = \frac{\sqrt{5}-1}{2}$ diese Gleichung. Das Risiko Y ist nach Definition identisch verteilt wie das Risiko X , deshalb gilt:

$$H(X) + H(Y) = 2H(X) = \sqrt{5} - 1.$$

Nun soll $H(X + Y)$ berechnet werden. Es gilt:

$$\begin{aligned} v(k) &= E(v(k + H(X + Y) - (X + Y))) \\ 0 &= E(\log(H(X + Y) - (X + Y) + 1)) \\ 0 &= \frac{1}{4} \log(H(X + Y) + 1) + \frac{1}{4} \log(H(X + Y) - 1) + \frac{1}{2} \log(H(X + Y)) \\ 0 &= \log(H(X + Y) + 1) + \log(H(X + Y) - 1) + 2 \log(H(X + Y)) \\ 0 &= \log((H(X + Y) + 1)(H(X + Y) - 1)(H(X + Y))^2) \\ 1 &= (H(X + Y) + 1)(H(X + Y) - 1)(H(X + Y))^2 \\ 1 &= (H(X + Y)^2 - 1)(H(X + Y))^2 \\ 0 &= (H(X + Y))^4 - (H(X + Y))^2 - 1. \end{aligned}$$

Aufgrund der Beschränkung auf positive Prämien wird die obige Gleichung nur von $H(X + Y) = \sqrt{\frac{1+\sqrt{5}}{2}}$ gelöst.

Insgesamt gilt daher:

$$H(X) + H(Y) = \sqrt{5} - 1 \neq \sqrt{\frac{1 + \sqrt{5}}{2}} = H(X + Y).$$

Positive Homogenität: Das oben gegebene Beispiel zur Additivität kann auch benutzt werden um zu zeigen, dass das Nullnutzenprinzip nicht positiv homogen ist. Für $H(2X)$ gilt:

$$\begin{aligned} v(k) &= E(v(k + H(2X) - 2X)) \\ 0 &= E(\log(H(2X) - 2X + 1)) \\ 0 &= \log(H(2X) + 1) + \log(H(2X) - 1) \\ 1 &= (H(2X) + 1)(H(2X) - 1) \\ 1 &= (H(2X))^2 - 1 \\ 2 &= (H(2X))^2. \end{aligned}$$

Nur der positive Wert $H(2X) = \sqrt{2}$ löst diese Gleichung. Es gilt daher:

$$H(2X) = \sqrt{2} \neq \sqrt{5} - 1 = 2H(X).$$

Damit ist gezeigt, dass das Nullnutzenprinzip nicht positiv homogen ist. $\square$

Bemerkung. Wählt man $v(x) = x$ als Nutzenfunktion dann erhält man für $H(X)$ das Nettoprämienprinzip.

Das Exponentialprinzip

Nach dem allgemeinen Nullnutzenprinzip wird nun ein spezielles Nullnutzenprinzip, das Exponentialprinzip, behandelt:

Definition 4.10.

Sei X ein versicherbares Risiko und sei die Nutzenfunktion v durch

$$v(x) := 1 - e^{-\beta x} \quad \beta > 0$$

definiert. Sei weiters k die Anfangsrisikoreserve. Dann wird das Prämienprinzip $H(X)$, das die Gleichung

$$v(k) = E(v(k + H(X) - X))$$

erfüllt, als Exponentialprinzip bezeichnet.

Satz 4.6. *Das Exponentialprinzip lässt sich explizit durch*

$$H(X) = \frac{1}{\beta} \log(E(e^{\beta X}))$$

darstellen.

Beweis.

Für $H(X)$ muss $v(k) = E(v(k + H(X) - X))$ mit $v(k) = 1 - e^{-\beta k}$ gelten. Für den Erwartungswert $E(v(k + H(X) - X))$ gilt:

$$\begin{aligned} E(v(k + H(X) - X)) &= E(1 - e^{-\beta(k+H(X)-X)}) \\ &= 1 - E(e^{-\beta k - \beta H(X) + \beta X}) = \\ &= 1 - E(e^{-\beta k} e^{-\beta H(X)} e^{\beta X}) = \\ &= 1 - e^{-\beta k} e^{-\beta H(X)} E(e^{\beta X}). \end{aligned}$$

Deshalb gilt:

$$\begin{aligned} v(k) &= E(v(k + H(X) - X)) \\ 1 - e^{-\beta k} &= 1 - e^{-\beta k} e^{-\beta H(X)} E(e^{\beta X}) \\ e^{-\beta k} &= e^{-\beta k} e^{-\beta H(X)} E(e^{\beta X}) \\ e^{\beta H(X)} &= E(e^{\beta X}) \\ \log(e^{\beta H(X)}) &= \log(E(e^{\beta X})) \\ \beta H(X) &= \log(E(e^{\beta X})) \\ H(X) &= \frac{1}{\beta} \log(E(e^{\beta X})). \end{aligned}$$

□

Satz 4.7. *Das Exponentialprinzip ist translationsinvariant, maximalschadenbegrenzt, additiv und nicht positiv homogen.*

Beweis.

Dass das Exponentialprinzip translationsinvariant und maximalschadenbegrenzt ist folgt daraus, dass diese Eigenschaften generell für ein Nullnutzenprinzip gelten.

Allerdings ist das Exponentialprinzip additiv, was für das Nullnutzenprinzip im Allgemeinen nicht gilt: Seien X und Y zwei unabhängige, versicherbare Risiken, dann gilt für

das Exponentialprinzip $H(X)$ mit dem Parameter $\beta > 0$:

$$\begin{aligned}
 H(X + Y) &= \frac{1}{\beta} \log(E(e^{\beta(X+Y)})) = \\
 &= \frac{1}{\beta} \log(E(e^{\beta X} e^{\beta Y})) = \\
 &= \frac{1}{\beta} \log(E(e^{\beta X}) E(e^{\beta Y})) = \\
 &= \frac{1}{\beta} (\log(E(e^{\beta X})) + \log(E(e^{\beta Y}))) = \\
 &= \frac{1}{\beta} \log(E(e^{\beta X})) + \frac{1}{\beta} \log(E(e^{\beta Y})) = H(X) + H(Y).
 \end{aligned}$$

Dass das Exponentialprinzip nicht positiv homogen ist kann mit folgendem Gegenbeispiel gezeigt werden: Sei das Risiko X folgendermaßen verteilt:

$$P(X = 0) = P(X = 1) = \frac{1}{2}.$$

Dann gilt:

$$P(2X = 0) = P(2X = 2) = \frac{1}{2}.$$

Für die Prämie $H(X)$ folgt daraus mit $\beta = 1$:

$$\begin{aligned}
 H(X) &= \frac{1}{\beta} \log(E(e^{\beta X})) = \log(E(e^X)) \\
 &= \log\left(\frac{1}{2}e^0 + \frac{1}{2}e^1\right) = \log\left(\frac{1+e}{2}\right).
 \end{aligned}$$

Für $H(2X)$ gilt dann:

$$\begin{aligned}
 H(2X) &= \frac{1}{\beta} \log(E(e^{\beta 2X})) = \log(E(e^2 X)) = \\
 &= \log\left(\frac{1}{2}e^0 + \frac{1}{2}e^2\right) = \log\left(\frac{1+e^2}{2}\right).
 \end{aligned}$$

Insgesamt gilt deshalb:

$$2H(X) = 2\log\left(\frac{1+e}{2}\right) \neq \log\left(\frac{1+e^2}{2}\right) = H(2X).$$

□

4.2.4 Das Esscherprinzip

Das Esscherprinzip ist ein Prämienprinzip, das ökonomische Überlegungen in die Prämienkalkulation einschließt. Diese Überlegungen werden nach der Definition des Prinzips näher erläutert.

Definition 4.11.

Sei X ein versicherbares Risiko und H ein Prämienprinzip, das durch

$$H(X) := \frac{E(Xe^{\beta X})}{E(e^{\beta X})}$$

mit $\beta \in \mathbb{R}, \beta > 0$ definiert ist. Dann wird H als Esscherprinzip bezeichnet.

Bemerkung. In der Literatur wird gelegentlich allgemeiner $H(X) := \frac{E(Xg(X))}{E(g(X))}$ mit $g : (0, \infty) \rightarrow (0, \infty)$, g monoton wachsend als Esscherprinzip verstanden, dies wird hier im Folgenden als allgemeines Esscherprinzip bezeichnet. Gilt im allgemeinen Esscherprinzip $g(x) = x^\alpha$ mit $\alpha \in \mathbb{N}$, dann wird es in der Literatur auch als Karlsruhe-Prinzip bezeichnet.

Aus dem allgemeinen Esscherprinzip lässt sich das Nettoprämienprinzip ableiten, wenn man $g(x) := 1$ setzt.

Das Esscherprinzip geht auf den Mathematiker Hans Bühlmann zurück, der es 1980 als ökonomisches Prämienprinzip vorschlug (vgl. [Sch97, S. 125]). Im Gegensatz zu den bisher dargestellten Prämienprinzipien wird bei diesem Prinzip nicht nur die eigene Bewertung eines Risikos in die Prämienberechnung einbezogen, sondern auch die Bewertung der Konkurrenz. Die Grundgedanken dieser Überlegungen sind folgende (nach [Sch97, S. 126] und [Büh88, S. 125f]): Es wird davon ausgegangen, dass am Markt n Versicherungsunternehmen tätig sind. Jedes Versicherungsunternehmen i ($i = 1, \dots, n$) bewertet ein Risiko X mit der Nutzenfunktion $v_i(x) = \frac{1}{a_i}(1 - e^{-a_i x})$. Außerdem sollen alle am Markt tätigen Unternehmen Risiken untereinander austauschen können, bis jedes Unternehmen ein Optimum erreicht, d. h. bis ein Pareto-Optimum eintritt. Die sogenannte Risikoaversion $-\frac{v_i''(x)}{v_i'(x)}$ ist in diesem Fall konstant und hat für das Versicherungsunternehmen i den Wert a_i . Bühlmann zeigte, dass es in diesem Fall (bei unabhängigen Risiken) genau einen Gleichgewichtspreis gibt, nämlich:

$$\frac{E(Xe^{\beta X})}{E(e^{\beta X})},$$

mit

$$\frac{1}{\beta} = \sum_{i=1}^n \frac{1}{a_i}.$$

Dies entspricht genau der Esscherprämie, der Parameter β gibt dabei ein Maß für die mittlere Risikoaversion des Marktes an.

Aufgrund der oben dargelegten Überlegungen, die zum Esscherprinzip führten, muss die Frage gestellt werden, wie realitätsnah dieses Prämienprinzip einen Marktpreis angeben kann. In [Sch97, S. 127] wird kritisiert, dass der beim Esscherprinzip ermittelte Wert alleine vom Markt, das heißt von der Bewertung der Mitbewerber abhängt. Als noch größeres Problem wird es gesehen, dass angenommen wird, dass die Versicherungsunternehmen beliebig Risiken untereinander austauschen können, was zudem kostenfrei für die einzelnen Unternehmen geschehen müsste. Dies ist in einer realen Marktsituation nur schwer denkbar. Außerdem wird angemerkt, dass sich dieses Prämienprinzip nur dann einsetzen lässt, wenn es mehrere Unternehmen gibt, die eine bestimmte Versicherung anbieten. Ist man der einzige Anbieter dieser Versicherung, dann lässt sich mit dem Esscherprinzip keine Prämie berechnen.

Durch die Verwendung der gleichen Nutzenfunktion ergibt sich auch eine Verbindung zwischen dem Esscherprinzip und dem Exponentialprinzip: Allerdings wird beim Exponentialprinzip die Prämie nur aus der Sicht des eigenen Unternehmens ermittelt. Beim Esscherprinzip hingegen wird ein Versicherungsmarkt betrachtet, auf dem alle tätigen Versicherungsunternehmen nach ihren eigenen Nutzenfunktionen bewerten und entscheiden (vgl. [4, S. 15]).

Im Folgenden wird das Esscherprinzip auf seine Eigenschaften hin untersucht:

Satz 4.8. *Das Esscherprinzip ist additiv, translationsinvariant, maximalschadenbegrenzt und nicht positiv homogen.*

Beweis. nach [11, S. 8f]

Seien X und Y zwei unabhängige, versicherbare Risiken. Dann gilt bezüglich der zu untersuchenden Eigenschaften:

Additivität:

$$\begin{aligned} H(X + Y) &= \frac{E((X + Y)e^{\beta(X+Y)})}{E(e^{\beta(X+Y)})} = \\ &= \frac{E(Xe^{\beta(X+Y)} + Ye^{\beta(X+Y)})}{E(e^{\beta(X+Y)})} = \end{aligned}$$

$$\begin{aligned}
&= \frac{E(Xe^{\beta X}e^{\beta Y} + Ye^{\beta X}e^{\beta Y}))}{E(e^{\beta X}e^{\beta Y})} = \\
&= \frac{E(Xe^{\beta X})E(e^{\beta Y}) + E(Ye^{\beta Y})E(e^{\beta X}))}{E(e^{\beta X})E(e^{\beta Y})} = \\
&= \frac{E(Xe^{\beta X})E(e^{\beta Y})}{E(e^{\beta X})E(e^{\beta Y})} + \frac{E(Ye^{\beta Y})E(e^{\beta X}))}{E(e^{\beta X})E(e^{\beta Y})} = \\
&= \frac{E(Xe^{\beta X})}{E(e^{\beta X})} + \frac{E(Ye^{\beta Y})}{E(e^{\beta Y})} = H(X) + H(Y).
\end{aligned}$$

Translationsinvarianz: Sei $a \in \mathbb{R}, a > 0$:

$$\begin{aligned}
H(X + a) &= \frac{E((X + a)e^{\beta(X+a)})}{E(e^{\beta(X+a)})} = \\
&= \frac{E(Xe^{\beta(X+a)} + ae^{\beta(X+a)})}{e^{\beta a}E(e^{\beta X})} = \\
&= \frac{E(Xe^{\beta X}e^{\beta a} + ae^{\beta X}e^{\beta a})}{e^{\beta a}E(e^{\beta X})} = \\
&= \frac{e^{\beta a}E(Xe^{\beta X}) + ae^{\beta a}E(e^{\beta X})}{e^{\beta a}E(e^{\beta X})} = \\
&= \frac{E(Xe^{\beta X})}{E(e^{\beta X})} + a = H(X) + a.
\end{aligned}$$

Maximalschadenbegrenztheit: Sei $x_m := \inf\{x \in \mathbb{R} : P(X \leq x) = 1\}$, dann gilt:

$$Xe^{\beta X} \leq x_me^{\beta X}.$$

Deshalb gilt auch:

$$H(X) = \frac{E(Xe^{\beta X})}{E(e^{\beta X})} \leq \frac{E(x_me^{\beta X})}{E(e^{\beta X})} = \frac{x_mE(e^{\beta X})}{E(e^{\beta X})} = x_m.$$

Positive Homogenität: Sei $a \in \mathbb{R}, a > 0$, dann gilt für $H(aX)$:

$$H(aX) = \frac{E(aXe^{\beta aX})}{E(e^{\beta aX})} = a \frac{E(Xe^{\beta aX})}{E(e^{\beta aX})}.$$

Somit gilt auch:

$$aH(X) = a \frac{E(Xe^{\beta X})}{E(e^{\beta X})} \neq a \frac{E(Xe^{\beta aX})}{E(e^{\beta aX})} = H(aX) \quad \forall a \neq 1.$$

Deshalb ist das Esscherprinzip nicht positiv homogen. $\square$

Bemerkung. Das allgemeine Esscherprinzip $H(X) := \frac{E(Xg(X))}{E(g(X))}$ mit $g : (0, \infty) \rightarrow (0, \infty)$, g monoton wachsend ist nur maximalschadenbegrenzt, aber nicht additiv, nicht translationsinvariant und nicht positiv homogen (vgl. [12, S. 103]).

4.2.5 Gewinnbeteiligung

In manchen Versicherungsverträgen werden die Prämien als schadenabhängig mit einer Gewinnbeteiligung für VersicherungsnehmerInnen festgelegt, das heißt, dass VersicherungsnehmerInnen am Beginn des Jahres eine Basisprämie bezahlen, je nach Höhe der Schäden wird am Ende des Jahres ein Teil dieser Basisprämie zurückerstattet. Dies lässt sich folgendermaßen modellieren (nach [8, S. 55]):

Sei X das zu versichernde Risiko, $H_0(X)$ die Basisprämie für dieses Risiko und Y eine Zufallsvariable mit $0 \leq Y \leq 1$ so, dass am Ende des Jahres $(1-Y)H_0(X)$ zurückerstattet wird. Insgesamt kann das Versicherungsunternehmen in diesem Fall Einnahmen von $YH_0(X)$ erwarten. Anfallende Zinsen werden hier nicht explizit berücksichtigt, die Zufallsvariable Y kann aber so angesetzt werden, dass eine implizite Berücksichtigung der Zinsen erfolgt. Die Basisprämie $H_0(X)$ muss dann so hoch sein, dass sie folgende Gleichungen erfüllt:

- Nettoprämienprinzip und Erwartungswertprinzip: $E(YH_0(X)) = E(X)$.
- Varianzprinzip: $E(YH_0(X)) + cV(YH_0(X)) = E(X) + cV(X)$.
- Nullnutzenprinzip: $v(k) = E(v(k + YH_0(X) - X))$.
- Exponentialprinzip: $\frac{1}{\beta} \log(E(e^{\beta X})) = \frac{1}{\beta} \log(E(e^{\beta YH_0(X)}))$. Das ist äquivalent mit: $E(e^{\beta X - \beta YH_0(X)}) = 1$.
- Esscherprinzip: $\frac{E(Xe^{\beta X})}{E(e^{\beta X})} = \frac{E(YH_0(X)e^{\beta YH_0(X)})}{E(e^{\beta YH_0(X)})}$.

Dass VersicherungsnehmerInnen am Gewinn des Unternehmens beteiligt werden kann verschiedene Gründe haben, beispielsweise weil es gesetzlich vorgeschrieben ist, weil andere Versicherungsunternehmen dies ebenfalls anbieten oder weil das Versicherungsunternehmen dadurch weniger bzw. geringere Schäden erwarten kann (Es kann statistisch nachgewiesen und psychologisch erklärt werden, dass es zwischen verschiedenen Arten von Versicherungsverträgen und der Schadeneintrittswahrscheinlichkeit einen Zusammenhang gibt, vgl. [7, S. 4]).

4.2.6 Übersicht und Fazit

In diesem Abschnitt wurden folgende Prämienprinzipien behandelt, die sich hinsichtlich ihrer Ansätze und Eigenschaften stark voneinander unterscheiden:

	additiv	translationsinvariant	pos. homogen	maximalschadenbegrenzt
Nettoprämienprinzip	✓	✓	✓	✓
Erwartungswertprinzip	✓	-	✓	-
Varianzprinzip	✓	✓	-	-
Standardabweichungsprinzip	-	✓	✓	-
allg. Nullnutzenprinzip	-	✓	-	✓
Exponentialprinzip	✓	✓	-	✓
Esscherprinzip	✓	✓	-	✓

Es ist in der obigen Tabelle ersichtlich, dass bis auf das Nettoprämienprinzip kein Prämienprinzip alle angeführten Eigenschaften erfüllt.

In der Literatur findet man oftmals ein weiteres Prämienprinzip, das Maximalschadenprinzip, das die Prämie als den höchstmöglichen Schaden ansetzt, also

$$H(X) := \inf\{x \in \mathbb{R} : P(X \leq x) = 1\}.$$

Der Grund, warum dieses Prämienprinzip immer wieder angeführt wird, ist die Tatsache, dass das Maximalschadenprinzip mit Ausnahme des Nettoprämienprinzips das einzige Prämienprinzip ist, das alle Eigenschaften erfüllt. Allerdings ist das Maximalschadenprinzip nur von theoretischem Interesse, für die Praxis ist es nicht relevant, da wohl keine Versicherungsnehmerin und kein Versicherungsnehmer bereit sein wird, eine Prämie in der Höhe des maximal möglichen Schadens zu bezahlen.

Hier wurde bei der Berechnung einer Prämie stets von einem einzelnen Risiko ausgegangen, für das die Prämie ermittelt werden soll. Praktisch aber wird bei der Prämienberechnung oft so vorgegangen, dass vorerst eine Prämie für den gesamten Bestand ermittelt wird und diese dann auf die einzelnen Risiken verteilt wird (vgl. [Sch02, S. 239]). Dabei kann es vorkommen, dass bei der Ermittlung der Prämie für den Bestand und bei der Verteilung auf die einzelnen Risiken unterschiedliche Prämienprinzipien zur Anwendung kommen, also dass beispielsweise die Prämie für den gesamten Bestand mit dem Nullnutzenprinzip berechnet wird und dass diese Prämie dann mit dem Varianzprinzip auf die einzelnen Risiken verteilt wird.

4.3 Erfahrungstarifizierung

4.3.1 Grundsätzliches zur Erfahrungstarifizierung

Unter der Erfahrungstarifizierung versteht man Konzepte, bei denen die Prämienhöhe teilweise oder ganz vom tatsächlichen Schadenverlauf eines Einzelrisikos abhängt (vgl. [HK91, S. 269]). Das heißt, dass die Prämie nicht durch gewisse Eigenschaften wie zum Beispiel dem Erwartungswert oder der Varianz eines Risikos kalkuliert wird, sondern dass die Prämie entsprechend der Anzahl an Schäden und evtl. deren Höhe in vergangenen Perioden festgelegt wird. Dabei wird die Tatsache berücksichtigt, dass Risiken, die sich a priori nicht unterscheiden (also z. B. die gleiche Verteilung besitzen) im Allgemeinen dennoch einen anderen Schadenverlauf aufweisen, wodurch unterschiedlich hohe Kosten für das Versicherungsunternehmen entstehen.

Durch die Erfahrungstarifizierung können folgende Effekte erzielt werden (nach [HK91, S. 270]):

- **Differenzierung der Prämie:** Die Prämien werden für ein konkretes Risiko anhand des individuellen Schadenverlaufs kalkuliert. Da die Prämien vom bisherigen Verlauf der Schäden abhängen nähern sie sich der individuellen Schadenerwartung an.
- **Überwälzung des Risikos:** Es verringert sich die Schwankung für das Versicherungsunternehmen, ein Teil des Risikos wird von den VersicherungsnehmerInnen getragen.
- **Selbstbeteiligung von VersicherungsnehmerInnen:** Da bereits aufgetretene Schäden in Zukunft zu einer Erhöhung der Prämie führen tragen VersicherungsnehmerInnen mit der höheren Zahlung an das Versicherungsunternehmen selbst zur Begleichung ihrer Schäden bei.

Erfahrungstarife können sowohl für ein Versicherungsunternehmen, als auch für VersicherungsnehmerInnen vorteilhaft sein, beispielsweise dann, wenn der erwartete Schaden nur ungenau bestimmt werden kann und deshalb ein hoher Schwankungszuschlag berechnet werden müsste. Außerdem haben VersicherungsnehmerInnen einen Nutzen, wenn sie (sowohl bewusst als auch unbewusst) dafür sorgen, dass Schäden vermieden werden. Eine Vermeidung von Schäden hat aber auch für das Versicherungsunternehmen positive Auswirkungen.

Auch im Hinblick auf Fairness bei der Prämienermittlung kann die Erfahrungstarifizierung begründet werden. So ist es vermutlich sowohl für Versicherungsunternehmen als auch für VersicherungsnehmerInnen einsichtig, dass für Versicherungsverträge, bei denen in der Vergangenheit nur wenige Schäden aufgetreten sind, niedrigere Prämien berechnet werden als bei Verträgen mit vielen Schäden.

Da die Prämie vom individuellen Schadenverlauf eines Risikos abhängt kann es bei der Erfahrungstarifizierung auch vorkommen, dass sich die Prämie von zwei Risiken mit identischer Verteilung unterscheidet, wenn der Schadenverlauf unterschiedlich ist. Bei der Prämienberechnung mit Prämienprinzipen könnte das nicht auftreten.

Die Berücksichtigung des Schadenverlaufs bei der Prämienbestimmung wird auch als sekundäre Prämiendifferenzierung bezeichnet (vgl. [Far06, S. 71]). (Unter der primären Prämiendifferenzierung versteht man das Zusammenfassen von Risiken mit gleichen Merkmalen in Beständen oder Teilbeständen.)

Die Methoden der Erfahrungstarifizierung sind sehr umfassend und in vielen Fällen sehr speziell. Im folgenden Abschnitt soll diese Art der Prämienbestimmung am Beispiel des Bonus-/Malus-Systems vorgestellt werden. Ein weiterer Ansatz, auf den hier aber nicht genauer eingegangen wird, ist die Credibility-Theorie, bei der die Prämie als gewichtetes Mittel aus der individuellen Schadenerfahrung eines Risikos und dem Schadenbedarf des gesamten Bestandes festgelegt wird (vgl. [Hei87, S. 138]).

Sinnvoll einsetzen lässt sich die Erfahrungstarifizierung vor allem in nicht homogenen Beständen und in Beständen mit einer großen Schadenhäufigkeit, praktisch findet sie vor allem in der KFZ-Haftpflichtversicherung, in der industriellen Feuerversicherung und in der Hagelversicherung Anwendung (vgl. [HK91, S. 269, S. 273] und [FW01, S. 226]).

4.3.2 Erfahrungstarife anhand von Bonus-/Malus-Systemen

Ein Bonus-/Malus-System (manchmal auch nur als Bonussystem bezeichnet) ist ein Modell zur Prämienfestlegung, in dem einzelne VersicherungsnehmerInnen einen Rabatt erhalten, wenn in einer Versicherungsperiode kein Schaden aufgetreten ist. Ist hingegen ein Schaden aufgetreten, dann muss die Versicherungsnehmerin oder der Versicherungsnehmer mit einem Zuschlag zur bisherigen Prämie rechnen. Der gewährte Rabatt bei Schadenfreiheit wird auch als Bonus bezeichnet, die Erhöhung der Prämie beim Eintreten von Schäden als Malus, wodurch dieses System seinen Namen erhalten hat.

Charakteristische Eigenschaften eines Bonus-/Malus-Systems

Ein Bonus-/Malus-System lässt sich allgemein folgendermaßen charakterisieren (nach [KS88, S. 92]):

- Die festgelegten Versicherungsperioden gelten für alle Risiken des Bestandes. Üblicherweise dauert eine Versicherungsperiode genau ein Jahr.
- Es gibt K Bonusklassen, die mit $1, 2, \dots, K$ durchnummeriert werden. Jedes Risiko ist in einer Versicherungsperiode genau einer Bonusklasse zugeordnet.
- Außerdem wird eine Prämienskala $\Pi = (H(1), H(2), \dots, H(K))$ definiert. Dabei bezeichnet $H(i)$ mit $1 \leq i \leq K$ die Prämie in der i -ten Bonusklasse. Für die einzelnen Prämien $H(i)$ gilt dabei per Konvention, dass $H(i) \leq H(i+1)$ für alle $1 \leq i < K$, das heißt in einer höheren Bonusklasse muss von VersicherungsnehmerInnen eine höhere oder eine zumindest gleich hohe Prämie bezahlt werden. (Die hier angeführte Konvention, dass in der höchsten Klasse die höchste Prämie und in der niedrigsten Klasse die geringste Prämie bezahlt werden muss wurde aufgrund des in Österreich gebräuchlichen Bonus-/Malus-Systems festgelegt. In anderen Staaten können anderen Konventionen üblich sein, beispielsweise ist im norwegischen Bonus-/Malus System die höchste Bonusklasse die Klasse mit der niedrigsten Prämie (vgl. [Sun99, S. 78]).)
- Es gibt eine Einstiegsklasse k_0 . Am Beginn eines Versicherungsvertrages wird jedes Risiko dieser Klasse zugewiesen.
- Für den Übergang von einer Bonusklasse in eine andere Bonusklasse werden Regeln definiert. Wie dieser Übergang erfolgt hängt von der individuellen Schadenerfahrung eines Risikos ab, üblicherweise wird dafür die Anzahl der Schäden in der betrachteten Periode herangezogen, es kann aber auch vorkommen, dass mehrere Perioden betrachtet werden, hier wird im folgenden aber nur auf die Betrachtung einer Periode eingegangen. Betrachtet man nur eine Periode, dann ist die Vorgangsweise folgende: War ein Risiko in der Bonusklasse i und hat r Schäden ($r \in \mathbb{N}$) verursacht, dann legt die Bonus-/Malus-Regel $T(i, r)$ fest, welcher Klasse dieses Risiko in der folgenden Periode zugeordnet wird. Beispielsweise würde für ein betrachtetes Risiko die Bonus-/Malus-Regel $T(3, 1) = 6$ bedeuten, dass ein Risiko in der Bonusklasse 3, das genau einen Schaden verursacht hat, in der nächsten Periode der Bonusklasse 6 zugeordnet wird. Für praktische Anwendungen gilt dabei, dass

eine höhere Zahl an Schadensfällen zur Folge hat, dass die zukünftig zugeordnete Bonusklasse höher oder zumindest gleich hoch ist, also dass $T(i, r + 1) \geq T(i, r)$ gilt.

Das Bonus-/Malus-System in der österreichischen KFZ-Haftpflichtversicherung

Ein Beispiel für die Anwendung einer Bonus-/Malus-Tarifierung ist die österreichische KFZ-Haftpflichtversicherung. Da der Abschluss einer KFZ-Haftpflichtversicherung notwendig ist um eine Zulassung für den Straßenverkehr zu erhalten ist sie für einen großen Teil der österreichischen Bevölkerung von Bedeutung. Die folgende Beschreibung erfolgt nach [14, S. 5ff]:

Das Bonus-/Malus-System im Bereich der KFZ-Haftpflichtversicherung wurde in Österreich im Jahr 1977 eingeführt und umfasst 18 Bonusklassen (die auch als Prämienstufen bezeichnet werden). Diese Bonusklassen werden mit 0, 1, ..., 17 durchnummeriert, wobei die Bonusklasse 0 die für einzelne VersicherungsnehmerInnen günstigste Stufe darstellt, die Bonusklasse 17 ist die für eine Versicherungsnehmerin oder einen Versicherungsnehmer ungünstigste Stufe. Der Einstieg bei Abschluss eines Versicherungsvertrages erfolgt in der Bonusklasse 9, die Grundprämie in dieser Bonusklasse wird je nach der Motorleistung des zu versichernden Fahrzeuges festgesetzt. Für die restlichen Bonusklassen erfolgt die Festsetzung der Prämie relativ zur Grundprämie, die Höhe der Prämie in den verschiedenen Prämienstufen ist in folgender Tabelle dargestellt:

Prämienstufe	Prämie (relativ zur Grundprämie)
0	50 %
1	50 %
2	60 %
3	60 %
4	70 %
5	70 %
6	80 %
7	80 %
8	100 %
9	100 %
10	120 %
11	120 %

12	140 %
13	140 %
14	170 %
15	170 %
16	200 %
17	200 %

Wie in der Tabelle ersichtlich umfassen die 18 Prämienstufen neun Tarifklassen, das heißt die Prämie steigt nicht nach jeder Stufe, sondern nach jeweils zwei Stufen an. Inzwischen sind Versicherungen allerdings nicht mehr gezwungen, sich starr an diese Festlegung der Prämien zu halten. So gewähren manche Versicherungsunternehmen Rabatte für bestimmte Bevölkerungsgruppen, beispielsweise Rabatte für Senioren.

Die Umstufung in diesem System erfolgt folgendermaßen:

- Hat eine Versicherungsnehmerin oder ein Versicherungsnehmer im Beobachtungszeitraum (in der Regel umfasst dieser ein Jahr beginnend mit dem 1. Oktober) keinen Schaden verursacht, dann wird sie oder er in der darauffolgenden Periode eine Prämienstufe tiefer eingestuft. Eine Ausnahme bildet hier die tiefste Prämienstufe, in dieser erfolgt keine Umstufung, falls kein Schaden verursacht wurde. Allerdings bieten viele Versicherungsunternehmen in der tiefsten Prämienstufe zusätzliche Boni an. Ein Beispiel dafür ist der sogenannte „Freischaden“, dabei führt ein verursachter Schaden nicht zu einer Umstufung, sondern es wird auch weiterhin die Prämie der tiefsten Stufe eingehoben.
- Für jeden verursachten Schaden erfolgt eine Erhöhung der Prämie um drei Prämienstufen, wobei die 17. Stufe die höchste Stufe bildet. Falls VersicherungsnehmerInnen eine schlechtere Einstufung und somit höhere Prämien abwenden wollen, dann haben sie die Möglichkeit, Schäden selbst zu begleichen, was sich besonders bei geringen Schäden lohnt. Der Effekt, dass VersicherungsnehmerInnen insbesondere kleine Schäden selbst begleichen, wird auch als „Bonushunger“ bezeichnet. Für ein Versicherungsunternehmen hat das den Vorteil, dass keine Kosten für die Schadensabwicklung entstehen.

Ermittlung einer optimalen Prämienskala

Nun soll eine Möglichkeit beschrieben werden, wie bei der Ermittlung der Prämien in den einzelnen Bonusstufen vorgegangen werden kann. Als Resultat erhält man eine Skala, bei der die Prämien in den jeweiligen Bonusstufen so gewählt werden, dass die Summe der mittleren quadratischen Abweichungen zwischen der Prämienhöhe eines Risikos und dem erwarteten Schadenaufwand für dieses Risiko minimiert wird. Die Beschreibung erfolgt nach [KS88, S. 92ff].

Im Folgenden wird die Prämie stets als Nettorisikoprämie verstanden, Sicherheitszuschläge werden hier nicht berücksichtigt.

Sei X ein Risiko und $N_X(n)$ die zugehörige Schadenzahl in der Periode n ($n = 0, 1, 2, \dots$). Die Zahl n gibt dabei an, wie viele Perioden das Risiko bereits versichert ist. Außerdem soll $Z_{X,T}(n)$ die Bonusklasse angeben, in der sich das Risiko X in der Periode n befindet, T bezeichnet dabei die Bonus-/Malus-Regel, die festlegt, wie die Umstufung am Ende der Periode erfolgt. $Z_{X,T}$ nimmt also stets einen Wert der Menge $\{1, 2, \dots, K\}$ an.

Am Beginn eines Versicherungsvertrages wird das Risiko X der Einstiegsklasse k_0 zugewiesen, das heißt für $n = 0$ gilt:

$$Z_{X,T}(0) = k_0.$$

Für $n > 0$ kann der Übergang mit der vom Zufall abhängigen Schadenzahl N_X des Risikos X folgendermaßen angegeben werden:

$$Z_{X,T}(n+1) = T(Z_{X,T}(n), N_X(n)).$$

Wie bereits erwähnt soll in diesem Modell die Summe der mittleren quadratischen Abweichungen zwischen dem erwarteten Schadenaufwand eines Risikos und der Prämienhöhe für dieses Risiko minimal sein. Der erwartete Schadenaufwand, also jener Betrag, der zur Begleichung von Schäden erwartet wird, wird für das Risiko X in der Periode n mit $m_X(n)$ bezeichnet. Eine optimale Prämienskala für die Periode n lässt sich dann folgendermaßen ermitteln: Seien $X_1, X_2, \dots, X_l$ alle in der Periode n versicherten Risiken. Die Prämienskala $\Pi = (H(1), H(2), \dots, H(K))$ in der Periode n muss so gewählt werden, dass der folgende Ausdruck ein Minimum annimmt:

$$\sum_{k=1}^l E \left((m_{X_k}(n) - H(Z_{X_k,T}(n)))^2 \right).$$

Mit dieser Vorgangsweise wird der Gedanke verfolgt, dass die Prämie in den jeweiligen Bonusstufen möglichst wenig vom erwarteten Schadenaufwand abweicht.

Auf jedes im Bestand vorhandene Risiko können in einer bestimmten Versicherungsperiode mehrere Einzelschäden fallen, die auch als Einzelschäden berücksichtigt werden müssen, da die Zahl der Schäden einer Periode zur Umstufung herangezogen wird. Betrachtet man ein Risiko X , dann gibt die Zufallsvariable N_X die Anzahl an Schadensfällen in einer Periode an. Bezeichnet man die Höhe der Einzelschäden in der betrachteten Periode mit $Y_1, Y_2, \dots$, dann lässt sich X als Zufallssumme anschreiben:

$$X = \sum_{k=1}^{N_X} Y_k.$$

Nimmt man an, dass sich die Verteilung der Einzelschäden und die Verteilung der individuellen Schadenzahl nicht ändern, dann ist der Erwartungswert von X konstant. Deshalb gilt dann für den erwarteten Schadenaufwand des Risikos X :

$$m_X(n) = E(X) = E\left(\sum_{k=1}^{N_X} Y_k\right) \quad \forall n \in \mathbb{N}.$$

Führt man wie im Kapitel über das kollektive Modell eine Zufallsvariable Y ein, die die typische Schadenshöhe angibt, dann gilt:

$$m_X(n) = E(N_X)E(Y) \quad \forall n \in \mathbb{N}.$$

Festlegung der Einstiegsklasse

Nachdem eine Prämienskala ermittelt wurde muss noch die Festlegung der Einstiegsklasse k_0 erfolgen. Die Einstiegsklasse k_0 soll so gewählt werden, dass der erwartete Schadenaufwand möglichst wenig von der Prämie abweicht. Hier wird wiederum die mittlere quadratische Abweichung herangezogen und die Einstiegsklasse k_0 wird so gewählt, dass folgendes gilt (vgl. [KS88, S. 93]):

$$E((m_X(n) - H(k_0))^2) \leq E((m_X(n) - H(k))^2) \quad \forall k \in \{1, 2, \dots, K\}.$$

Ein Problem bei dieser Festlegung der Prämienskala ist es, dass sie jede Versicherungsperiode neu berechnet werden muss. Da sich die Bestände aber laufend ändern kann das dazu führen, dass sich die Prämienhöhen in einer Prämienstufe von Jahr zu Jahr

unterscheiden. In der praktischen Anwendung könnte das bedeuten, dass die Prämien möglicherweise stark ansteigen, wodurch eine Einstufung in eine niedrigere Klasse nicht mehr reizvoll ist. Praktisch muss deswegen versucht werden, die Bonusklassen und die Regeln zur Umstufung so festzulegen, dass sie über lange Zeit stabil bleiben.

Andere Ansätze zur Ermittlung von Bonus-/Malus-Systemen erhält man außerdem aus der Credibility-Theorie (vgl. [KS88, S. 94]).

5 Maßnahmen zur Risikominderung

Versicherungsunternehmen erhalten fixe Geldbeträge dafür, damit sie für eine ungewisse Anzahl an Schäden aufkommen, deren Schadenshöhen im Vorhinein nicht bekannt sind. Dadurch werden VersicherungsnehmerInnen geschützt, Versicherungsunternehmen aber tragen das versicherungstechnische Risiko. Dieses kann zwar nicht eliminiert, aber gemindert werden. Dazu werden hier zwei Strategien behandelt, die in der Praxis häufig anzutreffen sind, nämlich einerseits, dass das Versicherungsunternehmen selbst eine Versicherung abschließt (Rückversicherung) und andererseits, dass VersicherungsnehmerInnen für Teile der aufgetretenden Schäden aufkommen müssen (Selbstbehalte).

5.1 Rückversicherung

Bei einer Rückversicherung schließt ein Versicherungsunternehmen selbst einen Versicherungsvertrag mit einem Rückversicherer ab, um Versicherungsschutz zu erhalten, sie lässt sich deshalb als Versicherung charakterisieren, die die von einem Versicherungsunternehmen übernommene Gefahr umfasst (vgl. [Sch98, S. 10]). Die Motivation und die Gründe, einen Rückversicherungsvertrag abzuschließen, können vielfältig sein, als Hauptgrund kann aber nach [Lie09, S. 50] die Teilung und die Reduktion des versicherungstechnischen Risikos gesehen werden. Je nach Rückversicherungsform kann dies beispielsweise den Schutz vor Großschäden oder den Schutz vor starken Schwankungen der Schadenshöhen umfassen, aber Rückversicherung bietet auch dann einen Versicherungsschutz, wenn Fehler bei der Schätzung des übernommenen Risikos (zum Beispiel eine falsche Annahme zur Verteilung von Schäden) gemacht wurden. Darüber hinaus kann der Abschluss eines Rückversicherungsvertrages für ein Versicherungsunternehmen sinnvoll sein, wenn das Eigenkapital im Verhältnis zum Gesamtvolumen der abgeschlossenen Versicherungsverträge als zu gering eingestuft wird (vgl. [Str88a, S. 7]). Weiters können Rückversicherer Dienstleistungen anbieten, von denen ein Versicherungsunternehmen profitieren kann, solche Dienstleistungen können beispielsweise das

Schätzen und Prüfen von Risiken, die Bereitstellung von umfangreichen Statistiken oder Beratungstätigkeiten sein (vgl. [Str88a, S. 10]).

Ist ein Rückversicherungsvertrag abgeschlossen, dann werden anfallende Schäden zwischen dem Erstversicherer und dem Rückversicherer aufgeteilt. Ob dabei alle anfallenden Schäden oder nur bestimmte Schäden (z. B. Schäden, die eine bestimmte Schadenshöhe übersteigen) aufgeteilt werden, hängt von der Form der Rückversicherung ab. Außerdem ist es von der Rückversicherungsform abhängig, was als Schaden betrachtet wird, so können einzelne Schäden, der jährliche Gesamtschaden einzelner Verträge, der Gesamtschaden aller Verträge sowie alle Schäden, die ein bestimmtes Ereignis betreffen (z. B. alle Schäden eines Hagelschauers), aufgeteilt werden. Je nachdem, wie die Aufteilung erfolgt, können bestimmte Arten der Rückversicherung auf eine Art der Modellierung beschränkt sein, werden beispielsweise einzelne Schäden aufgeteilt, so muss ein kollektives Modell zugrunde gelegt werden, in dem diese Einzelschäden erfasst sind.

Grundsätzlich können Rückversicherungsverträge unterschieden werden in Verträge, bei denen Schäden proportional aufgeteilt werden und Verträge, bei denen die Aufteilung nicht proportional erfolgt (vgl. [Sch05, S. 170]):

- **Proportionale Rückversicherung:** Bei der proportionalen Rückversicherung erfolgt die Aufteilung so, dass der Erstversicherer einen gewissen prozentualen Teil des Schadens übernimmt, der Rückversicherer übernimmt den Rest des Schadens.
- **Nichtproportionale Rückversicherung:** Im Gegensatz zur proportionalen Rückversicherung übernimmt der Rückversicherer bei der nichtproportionalen Rückversicherung nur einen Anteil, wenn der Schaden eine vorher festgelegte Schadenshöhe übersteigt.

Ein wesentlicher Unterschied zwischen der proportionalen und der nichtproportionalen Rückversicherung ist die Beteiligung des Rückversicherers: Bei der proportionalen Rückversicherung ist dieser an jedem Schaden beteiligt, bei der nichtproportionalen Rückversicherung hingegen nur an gewissen Schäden.

Für VersicherungsnehmerInnen ist es im Allgemeinen nicht ersichtlich, ob ein Versicherungsunternehmen rückversichert ist, denn der Versicherungsvertrag wird nur zwischen ihnen selbst und dem Versicherungsunternehmen abgeschlossen. Der Rückversicherer hat dadurch auch keine direkten Verpflichtungen gegenüber den VersicherungsnehmerInnen des Erstversicherers, sämtliche vertragliche Verpflichtungen (wie beispielsweise die Schadensregulierung im Schadensfall) betreffen ausschließlich den Erstversicherer.

Die Funktionsweise der Rückversicherung ist identisch zu jener der Erstversicherung, auch bei der Rückversicherung soll ein Ausgleich im Kollektiv erreicht werden. Unternehmen, die Rückversicherungen anbieten, müssen nicht auf diese spezialisiert sein, sondern können auch als Erstversicherer auftreten, praktisch entspricht die Rückversicherung deshalb oft einer Aufteilung des Risikos auf mehrere Unternehmen (vgl. [Sch02, S. 183]).

Im Folgenden werden einige Arten der Rückversicherung behandelt, anschließend wird auf die Einsatzbereiche der verschiedenen Rückversicherungsarten eingegangen.

5.1.1 Proportionale Rückversicherung

In der proportionalen Rückversicherung werden anfallende Schäden in einem festen Verhältnis zwischen dem Erstversicherer und dem Rückversicherer aufgeteilt. Mathematisch lässt sich diese Form der Rückversicherung folgendermaßen charakterisieren (vgl. [Sch02, S. 184]): Von einem Schaden S (dies kann ein Einzelschaden, die Summe der jährlichen Schäden oder der jährliche Gesamtschaden eines Bestandes sein) übernimmt der Rückversicherer den Anteil qS mit $q \in (0, 1)$, der Erstversicherer übernimmt den Anteil $(1 - q)S$.

Quotenrückversicherung

Bei der Quotenrückversicherung wird ein prozentualer Anteil festgelegt, den ein Rückversicherer von jedem aufgetretenen Schaden übernimmt. Der vom Rückversicherer übernommene Anteil wird als Quote q bezeichnet und muss sinnvollerweise mehr als 0 % und weniger als 100 % betragen, das heißt, es muss $q \in (0, 1)$ gelten. Da alle anfallenden Schäden im gleichen Verhältnis aufgeteilt werden, bedeutet das für den Gesamtschaden, dass auch dieser in einem festgelegten Verhältnis zwischen dem Erstversicherer und dem Rückversicherer aufgeteilt wird (vgl. [Wol88, S. 298]).

Da es bei der Quotenrückversicherung nicht von Bedeutung ist, ob die Quote auf einzelne Schäden oder den jährlichen Gesamtschaden eines Versicherungsvertrages bezogen wird, kann dieser Art der Rückversicherung sowohl das individuelle, als auch das kollektive Modell zugrunde gelegt werden.

Betrachtet man den Gesamtschaden S , dann lässt sich der Anteil S_R , den der Rückversicherer übernimmt, folgendermaßen angeben:

$$S_R = qS.$$

Der Anteil S_E , für den der Erstversicherer aufkommt, beträgt deshalb:

$$S_E = (1 - q)S.$$

Individuelles Modell: Im individuellen Modell gilt für einen Bestand $\{X_1, X_2, \dots, X_n\}$ von n unabhängigen Risiken:

$$S_R = q \sum_{i=1}^n X_i,$$

$$S_E = (1 - q) \sum_{i=1}^n X_i.$$

Für die Erwartungswerte von S_R und S_E heißt das:

$$E(S_R) = E\left(q \sum_{i=1}^n X_i\right) = q \sum_{i=1}^n E(X_i) = qE(S),$$

$$E(S_E) = (1 - q) \sum_{i=1}^n E(X_i) = (1 - q)E(S).$$

Außerdem gilt für die Varianzen $V(S_R)$ und $V(S_E)$:

$$V(S_R) = V\left(q \sum_{i=1}^n X_i\right) = q^2 V\left(\sum_{i=1}^n X_i\right) = q^2 \sum_{i=1}^n V(X_i) = q^2 V(S),$$

$$V(S_E) = (1 - q)^2 \sum_{i=1}^n V(X_i) = (1 - q)^2 V(S).$$

Kollektives Modell: Im kollektiven Modell mit der Schadenzahl N und den identisch verteilten, unabhängigen Schadenshöhen $X_1, X_2, \dots$ mit der typischen Schadenshöhe $X := X_1$ gilt für die Schadenzahlen N_R und N_E von Rück- und Erstversicherer aufgrund der Tatsache, dass für beide gleich viele Schadensfälle auftreten:

$$N_R = N_E = N.$$

Deshalb gilt für die jeweiligen Gesamtschäden:

$$S_R = q \sum_{i=1}^N X_i,$$

$$S_E = (1 - q) \sum_{i=1}^N X_i.$$

Für die Erwartungswerte von S_R und S_E gilt:

$$\begin{aligned} E(S_R) &= E\left(q \sum_{i=1}^N X_i\right) = qE\left(\sum_{i=1}^N X_i\right) = qE(N)E(X) = qE(S), \\ E(S_E) &= (1 - q)E\left(\sum_{i=1}^N X_i\right) = (1 - q)E(N)E(X) = (1 - q)E(S). \end{aligned}$$

Außerdem gilt für $V(S_R)$ und $V(S_E)$:

$$\begin{aligned} V(S_R) &= V\left(q \sum_{i=1}^N X_i\right) = q^2 V\left(\sum_{i=1}^N X_i\right) = q^2 (E(N)V(X) + V(N)(E(X))^2) = q^2 V(S), \\ V(S_E) &= (1 - q)^2 V\left(\sum_{i=1}^N X_i\right) = (1 - q)^2 (E(N)V(X) + V(N)(E(X))^2) = (1 - q)^2 V(S). \end{aligned}$$

Summenexzedentenrückversicherung

Im Gegensatz zur Quotenrückversicherung werden bei der Summenexzedentenrückversicherung nur gewisse Risiken rückversichert, nämlich jene Risiken, deren Versicherungssumme einen festgelegten Betrag übersteigt. Dieser Betrag wird als Maximum $m \in \mathbb{R}$, $m > 0$ bezeichnet (vgl. [Sch02, S. 185]). Unter der Versicherungssumme wird dabei jener Betrag verstanden, den das Versicherungsunternehmen im Schadensfall höchstens für ein konkretes Risiko ausbezahlt.

Die Quote q gibt hier analog zur Quotenrückversicherung jenen Anteil an, den der Rückversicherer übernimmt, allerdings wird die Quote bei der Summenexzedentenrückversicherung in Abhängigkeit von der Versicherungssumme v festgelegt (vgl. [Sch02, S. 185]):

$$q(v) = \begin{cases} 0 & v \leq m, \\ 1 - \frac{m}{v} & v > m. \end{cases}$$

Im Falle eines Schadens wird hier folgendermaßen zwischen dem Erstversicherer und dem Rückversicherer aufgeteilt: Ist die Versicherungssumme des Risikos kleiner oder gleich dem Maximum m , dann trägt der Erstversicherer alle anfallenden Schäden vollständig. (In der praktischen Anwendung wird für solche Risiken keine Rückversicherung abge-

geschlossen, für die mathematische Betrachtung ist es aber einfacher, die Quote als $q = 0$ festzulegen.) Hat das Risiko eine höhere Versicherungssumme als m , dann übernimmt der Erstversicherer den Anteil $\frac{m}{v}$ und der Rückversicherer den Anteil $1 - \frac{m}{v}$.

Die Aufteilung bei der Summenexzedentenrückversicherung erfolgt also analog zur Quotenrückversicherung für ein konkretes Risiko in einem festen Verhältnis, allerdings werden nicht alle Schäden aufgeteilt, sondern nur die Schäden jener Risiken, deren Versicherungssumme einen bestimmten Wert übersteigt.

Bei dieser Art der Rückversicherung ist es notwendig zu wissen, welchem Vertrag ein angefallener Schaden zugeordnet werden kann. Da dies im kollektiven Modell nicht erfasst wird muss für die Summenexzedentenrückversicherung das individuelle Modell zugrunde gelegt werden.

Sei $\{X_1, X_2, \dots, X_n\}$ ein Bestand von n unabhängigen Risiken mit der Versicherungssumme v_i für das Risiko X_i ($i = 1, 2, \dots, n$). Dann gilt im individuellen Modell:

$$S_R = \sum_{i=1}^n q(v_i) X_i,$$

$$S_E = \sum_{i=1}^n (1 - q(v_i)) X_i.$$

Für die Erwartungswerte von S_R und S_E gilt:

$$E(S_R) = E\left(\sum_{i=1}^n q(v_i) X_i\right) = \sum_{i=1}^n E(q(v_i) X_i) = \sum_{i=1}^n q(v_i) E(X_i),$$

$$E(S_E) = E\left(\sum_{i=1}^n (1 - q(v_i)) X_i\right) = \sum_{i=1}^n E((1 - q(v_i)) X_i) = \sum_{i=1}^n (1 - q(v_i)) E(X_i).$$

Außerdem gilt für $V(S_R)$ und $V(S_E)$:

$$V(S_R) = V\left(\sum_{i=1}^n q(v_i) X_i\right) = \sum_{i=1}^n V(q(v_i) X_i) = \sum_{i=1}^n (q(v_i))^2 V(X_i),$$

$$V(S_E) = V\left(\sum_{i=1}^n (1 - q(v_i)) X_i\right) = \sum_{i=1}^n V((1 - q(v_i)) X_i) = \sum_{i=1}^n (1 - q(v_i))^2 V(X_i).$$

In der praktischen Anwendung der Summenexzedentenrückversicherung ist es üblich, die Haftung des Rückversicherers zu begrenzen. Legt man das k -fache des Maximums m mit $k \in \mathbb{N}, k \geq 1$ als obere Grenze für die Haftung des Rückversicherers fest, dann

ergibt sich für q (vgl. [Sch02, S. 187]):

$$q(v) = \begin{cases} 0 & v \leq m, \\ 1 - \frac{m}{v} & m < v \leq (k+1)m, \\ k\frac{m}{v} & v > (k+1)m. \end{cases}$$

Konkret heißt das also, dass der Erstversicherer voll für Risiken haftet, deren Versicherungssumme kleiner als das Maximum m ist. Bei Risiken mit einer Versicherungssumme zwischen m und $(k+1)m$ haftet der Erstversicherer mit dem Anteil $\frac{m}{v}$, der Rückversicherer mit dem Anteil $1 - \frac{m}{v}$. Hat ein Risiko eine höhere Versicherungssumme als $(k+1)m$, dann haftet der Rückversicherer mit dem Anteil $k\frac{m}{v}$, der Erstversicherer mit dem Anteil $1 - k\frac{m}{v}$.

Als Funktion betrachtet ist $q(v)$ monoton wachsend im Intervall $(m, (k+1)m]$ und monoton fallend im Intervall $((k+1)m, \infty)$, außerdem ist $q(v)$ stetig im Intervall (m, ∞) , deshalb nimmt $q(v)$ an der Stelle $(k+1)m$ ein Maximum an. Der Wert von $q(v)$ an dieser Stelle beträgt $\frac{k}{k+1}$, deshalb ist der höchste Betrag, den ein Rückversicherer bei dieser Beschränkung der Haftung für ein Risiko zu bezahlen hat km (dieser ergibt sich aus dem Produkt der Versicherungssumme $(k+1)m$ und der Quote $q((k+1)m) = \frac{k}{k+1}$; vgl. [Sch02, S. 187]).

5.1.2 Nichtproportionale Rückversicherung

Bei der nichtproportionalen Rückversicherung übernimmt der Rückversicherer nur einen Anteil des Schadens, wenn dieser eine bestimmte Höhe übersteigt. Mathematisch heißt das: Den Schaden S übernimmt der Erstversicherer bis zur Schadenshöhe $d \in \mathbb{R}, d > 0$ vollständig, alles was über die Schadenshöhe d hinausgeht, wird vom Rückversicherer übernommen (vgl. [Lie09, S. 167]). Der Anteil des Erstversicherers beträgt deshalb $\min\{S, d\}$, der Anteil des Rückversicherers beträgt $\max\{S - d, 0\}$.

Schadenexzedentenrückversicherung

In der Schadenexzedentenrückversicherung werden einzelne Schäden betrachtet. Zwischen dem Erstversicherer und dem Rückversicherer wird ein Betrag $d \in \mathbb{R}, d > 0$ (als Priorität bezeichnet) festgelegt, bis zu dem der Erstversicherer den Schaden vollständig begleicht, ist die Schadenshöhe größer als die Priorität, dann übernimmt der Rückversicherer jenen Betrag, der die Priorität übersteigt (vgl. [Sch02, S. 191]). Mit der Scha-

den Höhe X und der Priorität d gilt deshalb für den Anteil X_E , den der Erstversicherer übernimmt:

$$X_E = \min\{X, d\}.$$

Der Rückversicherer übernimmt den Anteil X_R :

$$X_R = \max\{X - d, 0\}.$$

Legt man das kollektive Modell mit den Schadenshöhen $X_1, X_2, \dots$ und der Schaden-
zahl N zugrunde, dann bedeutet das für die Gesamtschäden S_E und S_R von Erst- und
Rückversicherer:

$$S_E = \sum_{i=1}^N \min\{X_i, d\},$$

$$S_R = \sum_{i=1}^N \max\{X_i - d, 0\}.$$

Im Folgenden wird beschrieben, wie die Schadenzen, die Schadenshöhen und die
Gesamtschäden von Erst- und Rückversicherer modelliert werden können. Die Model-
lierung kann auf mehrere Arten erfolgen, die hier beschriebene Art der Modellierung
orientiert sich an den in [Mac02, S. 332ff], [10, S. 81ff] und [Sch02, 196ff] beschriebenen
Modellen. Am Ende des Abschnitts wird kurz ein alternatives Modell für die Scha-
denexzedentenrückversicherung angeführt und diskutiert. Im gesamten Abschnitt wird
ein kollektives Modell mit der Schadenzen N , den identisch verteilten, unabhängigen
Schadenshöhen $X_1, X_2, \dots$ und der typischen Schadenshöhe $X := X_1$ betrachtet.

Die Schadenzen des Erstversicherers

Da der Erstversicherer an allen auftretenden Schäden beteiligt ist, kann die Anzahl an
Schäden für ihn durch die Schadenzen N angegeben werden.

Die Schadenzen des Rückversicherers

Der Rückversicherer ist hingegen nur an Schäden beteiligt, deren Schadenshöhe größer
als die Priorität ist, deswegen gilt für die Schadenzen N_R des Rückversicherers: $N_R \leq N$.
Die Schadenzen des Rückversicherers kann folgendermaßen angegeben werden (vgl. [10,
S. 83]):

Sei I_i eine Zufallsvariable, die angibt, ob die Schadenshöhe X_i größer als die Priorität

ist:

$$I_i := \begin{cases} 0 & X_i \leq d, \\ 1 & X_i > d. \end{cases}$$

Außerdem wird zur Vereinfachung der Rechnung die Zufallsvariable I eingeführt, die sich auf die typische Schadenshöhe X bezieht:

$$I := \begin{cases} 0 & X \leq d, \\ 1 & X > d. \end{cases}$$

Für N_R gilt dann:

$$N_R = \sum_{i=1}^N I_i.$$

Der Erwartungswert und die Varianz von N_R lassen sich nun wie der Erwartungswert und die Varianz des Gesamtschadens im kollektiven Modell folgendermaßen angeben:

$$\begin{aligned} E(N_R) &= E(N)E(I), \\ V(N_R) &= E(N)V(I) + V(N)(E(I))^2. \end{aligned}$$

Für $E(I)$ und $V(I)$ gilt:

$$\begin{aligned} E(I) &= 0 \cdot P(X \leq d) + 1 \cdot P(X > d) = P(X > d), \\ V(I) &= E(I^2) - (E(I))^2 = 0 \cdot P(X \leq d) + 1 \cdot P(X > d) - (P(X > d))^2 = \\ &= P(X > d) - (P(X > d))^2. \end{aligned}$$

Definiert man $p_d := P(X > d)$, dann ergibt sich daraus für $E(N_R)$ und $V(N_R)$:

$$\begin{aligned} E(N_R) &= E(N)p_d, \\ V(N_R) &= E(N)(p_d - p_d^2) + V(N)p_d^2. \end{aligned}$$

Für die Verteilung von N_R gilt (vgl. [10, S. 83]):

$$P(N_R = k) = \sum_{n=0}^{\infty} P(N = n, N_R = k).$$

Mit der Formel der totalen Wahrscheinlichkeit ergibt sich daraus:

$$P(N_R = k) = \sum_{n=0}^{\infty} P(N_R = k | N = n) P(N = n).$$

Unter der Bedingung, dass insgesamt n Schäden aufgetreten sind (also $N = n$) gilt für die Wahrscheinlichkeit, dass $N_R = k$ ist:

$$P(N_R = k | N = n) = \binom{n}{k} p_d^k (1 - p_d)^{n-k}.$$

Dies gilt, da die Wahrscheinlichkeit für den Eintritt eines Schadens, dessen Schadenshöhe größer als die Priorität ist, durch p_d gegeben ist und weil die einzelnen Schadenshöhen als unabhängig angenommen werden. Setzt man $P(N_R = k | N = n)$ ein, dann ergibt sich für $P(N_R = k)$:

$$P(N_R = k) = \sum_{n=0}^{\infty} \binom{n}{k} p_d^k (1 - p_d)^{n-k} P(N = n).$$

Ist $n < k$, dann gilt:

$$\binom{n}{k} p_d^k (1 - p_d)^{n-k} = 0.$$

Deshalb kann die Summation bei $n = k$ begonnen werden und es gilt für die Verteilung von N_R :

$$P(N_R = k) = \sum_{n=k}^{\infty} \binom{n}{k} p_d^k (1 - p_d)^{n-k} P(N = n).$$

Nun soll die Verteilung von N_R untersucht werden, wenn man für die Schadenzahl N eine im Kapitel über das kollektive Modell behandelte Verteilung voraussetzt:

Satz 5.1. *Seien $X_1, X_2, \dots$ die unabhängigen und identisch verteilten Schadenshöhen im kollektiven Modell mit der typischen Schadenshöhe $X := X_1$ und der Schadenzahl N . Dann gilt für die Schadenzahl N_R des Rückversicherers bei der Schadenerzedentenrückversicherung:*

1. *Ist N binomialverteilt mit den Parametern m und p , dann ist N_R binomialverteilt mit den Parametern m und pp_d .*

2. Ist N poissonverteilt mit dem Parameter λ , dann ist N_R poissonverteilt mit dem Parameter λp_d .
3. Ist N negativ-binomialverteilt mit den Parametern r und p , dann ist N_R negativ-binomialverteilt mit den Parametern r und $\frac{p}{p+p_d-pp_d}$.

Beweis.

1. Sei N binomialverteilt mit den Parametern m und p , dann gilt für die Verteilung von N :

$$P(N = n) = \binom{m}{n} p^n (1-p)^{m-n}.$$

Deshalb gilt für die Verteilung von N_R :

$$\begin{aligned}
P(N_R = k) &= \sum_{n=k}^m \binom{n}{k} p_d^k (1-p_d)^{n-k} P(N = n) = \\
&= \sum_{n=k}^m \binom{n}{k} p_d^k (1-p_d)^{n-k} \binom{m}{n} p^n (1-p)^{m-n} = \\
&= p_d^k \sum_{n=k}^m \frac{n!}{k!(n-k)!} (1-p_d)^{n-k} \frac{m!}{n!(m-n)!} p^n (1-p)^{m-n} = \\
&= p_d^k p^k (1-p)^{m-k} \sum_{n=k}^m \frac{m!}{k!(n-k)!(m-n)!} (1-p_d)^{n-k} p^{n-k} (1-p)^{k-n} = \\
&= p_d^k p^k (1-p)^{m-k} \frac{m!}{k!} \underbrace{\sum_{n=k}^m \frac{1}{(n-k)!(m-n)!} \left(\frac{(1-p_d)p}{1-p} \right)^{k-n}}_{= \frac{1}{(m-k)!} \left(1 + \frac{(1-p_d)p}{1-p} \right)^{m-k} = \frac{1}{(m-k)!} \left(\frac{1-pp_d}{1-p} \right)^{m-k}} = \\
&= \underbrace{\frac{m!}{k!(m-k)!}}_{= \binom{m}{k}} p_d^k p^k (1-p)^{m-k} \left(\frac{1-pp_d}{1-p} \right)^{m-k} = \\
&= \binom{m}{k} (pp_d)^k (1-pp_d)^{m-k}.
\end{aligned}$$

N_R ist also binomialverteilt mit den Parametern m und pp_d .

2. Sei N poissonverteilt mit dem Parameter λ , dann gilt für die Verteilung von N :

$$P(N = n) = \frac{\lambda^n}{n!} e^{-\lambda}.$$

Deshalb gilt für die Verteilung von N_R :

$$\begin{aligned}
P(N_R = k) &= \sum_{n=k}^{\infty} \binom{n}{k} p_d^k (1 - p_d)^{n-k} P(N = n) = \\
&= \sum_{n=k}^{\infty} \binom{n}{k} p_d^k (1 - p_d)^{n-k} \frac{\lambda^n}{n!} e^{-\lambda} = \\
&= p_d^k e^{-\lambda} \sum_{n=k}^{\infty} \frac{n!}{k!(n-k)!} (1 - p_d)^{n-k} \frac{\lambda^n}{n!} = \\
&= \frac{p_d^k e^{-\lambda} \lambda^k}{k!} \underbrace{\sum_{n=k}^{\infty} \frac{(1 - p_d)^{n-k} \lambda^{n-k}}{(n-k)!}}_{=e^{(1-p_d)\lambda} = e^{\lambda} e^{-p_d \lambda}} = \\
&= \frac{p_d^k \lambda^k}{k!} e^{-p_d \lambda}.
\end{aligned}$$

N_R ist also poissonverteilt mit dem Parameter λp_d .

3. Sei N negativ-binomialverteilt mit den Parametern r und p , dann gilt für die Verteilung von N :

$$P(N = n) = \binom{r+n-1}{n} p^r (1-p)^n.$$

Deshalb gilt für die Verteilung von N_R :

$$\begin{aligned}
P(N_R = k) &= \sum_{n=k}^{\infty} \binom{n}{k} p_d^k (1 - p_d)^{n-k} P(N = n) = \\
&= \sum_{n=k}^{\infty} \binom{n}{k} p_d^k (1 - p_d)^{n-k} \binom{r+n-1}{n} p^r (1-p)^n = \\
&= p_d^k p^r \sum_{n=k}^{\infty} \binom{n}{k} \binom{r+n-1}{n} (1 - p_d)^{n-k} (1-p)^n = \\
&= p_d^k p^r (1-p)^k \sum_{n=k}^{\infty} \frac{(r+n-1)!}{k!(n-k)!(r-1)!} (1 - p_d)^{n-k} (1-p)^{n-k} = \\
&= p_d^k p^r (1-p)^k \frac{1}{k!} \underbrace{\sum_{n=k}^{\infty} \frac{(r+n-1)!}{(n-k)!(r-1)!} ((1-p_d)(1-p))^{n-k}}_{= \frac{(r+k-1)!}{(r-1)!} (1-(1-p-p_d+p_d p))^{-r-k} = \frac{(r+k-1)!}{(r-1)!} (p+p_d-p_d p)^{-r-k}} =
\end{aligned}$$

$$\begin{aligned}
&= p_d^k p^r (1-p)^k \underbrace{\frac{1}{k!} \frac{(r+k-1)!}{(r-1)!}}_{=\binom{r+k-1}{k}} \frac{1}{(p+p_d-pp_d)^{r+k}} = \\
&= \binom{r+k-1}{k} \left(\frac{p}{p+p_d-pp_d} \right)^r \underbrace{\frac{(p_d(1-p))^k}{(p+p_d-pp_d)^k}}_{=\left(1-\frac{p}{p+p_d-pp_d}\right)^k} = \\
&= \binom{r+k-1}{k} \left(\frac{p}{p+p_d-pp_d} \right)^r \left(1 - \frac{p}{p+p_d-pp_d} \right)^k.
\end{aligned}$$

N_R ist also negativ-binomialverteilt mit den Parametern r und $\frac{p}{p+p_d-pp_d}$. $\square$

Bemerkung. Aufgrund dieses Satzes könnte man annehmen, dass N_R immer die gleiche Verteilung wie N besitzt. Allgemein gilt das aber nicht, ein Gegenbeispiel ist die logarithmische Verteilung: N_R ist nicht logarithmisch verteilt wenn N logarithmisch verteilt ist (vgl. [Mac02, S. 334]).

Aus der Verteilung von N ergibt sich nach dem obigen Satz für den Erwartungswert und die Varianz von N_R :

- Ist N binomialverteilt mit den Parametern m und p , dann gilt:

$$E(N_R) = mpp_d,$$

$$V(N_R) = mpp_d(1-pp_d).$$

- Ist N poissonverteilt mit dem Parameter λ , dann gilt:

$$E(N_R) = \lambda p_d,$$

$$V(N_R) = \lambda p_d.$$

- Ist N negativ-binomialverteilt mit den Parametern r und p , dann gilt:

$$E(N) = \frac{r(1 - \frac{p}{p+p_d-pp_d})}{\frac{p}{p+p_d-pp_d}} = \frac{r \frac{p_d-pp_d}{p+p_d-pp_d}}{\frac{p}{p+p_d-pp_d}} = \frac{rp_d(1-p)}{p},$$

$$V(N) = \frac{r(1 - \frac{p}{p+p_d-pp_d})}{\left(\frac{p}{p+p_d-pp_d}\right)^2} = \frac{r \frac{p_d-pp_d}{p+p_d-pp_d}}{\left(\frac{p}{p+p_d-pp_d}\right)^2} = \frac{rp_d(1-p)(p+p_d-pp_d)}{p^2}.$$

Die Schadenshöhen des Erstversicherers

Von einem Schaden der Höhe X übernimmt der Erstversicherer den Anteil $X_E = \max\{X, d\}$. Bezeichnet man die Verteilungsfunktion von X mit F , dann lässt sich die Verteilungsfunktion von X_E folgendermaßen angeben (vgl. [Mac02, S. 335]):

$$F_E(x) = P(X_E < x) = \begin{cases} F(x) & x < d, \\ 1 & x \geq d. \end{cases}$$

Im Allgemeinen ist das Auftreten von Schäden möglich, deren Schadenshöhe größer als die Priorität d ist (sonst würde der Abschluss eines Rückversicherungsvertrages keinen Sinn ergeben), deshalb gilt allgemein für die Verteilungsfunktion F , dass $F(d) \neq 1$. F_E ist somit im Allgemeinen nicht stetig im Punkt d und besitzt allgemein auch keine stetige Dichte. (Allerdings ist F_E , wie für eine Verteilungsfunktion gefordert, rechtsseitig stetig im Punkt d .)

Falls die Verteilungsfunktion F eine Dichte besitzt (bzw. falls die Dichte bekannt ist), dann lässt sich der Erwartungswert von X_E folgendermaßen berechnen (f bezeichnet dabei die Dichte von F):

$$\begin{aligned} E(X_E) &= \int_0^\infty x dF_E(x) = \\ &= \int_0^d x f(x) dx + \underbrace{\int_d^\infty d \cdot f(x) dx}_{=d \cdot \int_d^\infty f(x) dx = d \cdot (1 - F(d))} = \\ &= \int_0^d x f(x) dx + d \cdot (1 - F(d)). \end{aligned}$$

Das Integral $\int_0^d x f(x) dx$ kann durch partielle Integration mit $u(x) = x$ und $v' = f(x)$ auch folgendermaßen dargestellt werden:

$$\begin{aligned} \int_0^d x f(x) dx &= x F(x) \Big|_0^d - \int_0^d F(x) dx = \\ &= d \cdot F(d) - \int_0^d F(x) dx. \end{aligned}$$

Deshalb lässt sich der Erwartungswert von X_E auch folgendermaßen angeben:

$$\begin{aligned} E(X_E) &= \int_0^d x f(x) dx + d \cdot (1 - F(d)) = \\ &= d \cdot F(d) - \int_0^d F(x) dx + d \cdot (1 - F(d)) = \\ &= d - \int_0^d F(x) dx. \end{aligned}$$

(Obwohl zur Berechnung die Dichte von F verwendet wurde kann $E(X_E)$ auch so angegeben werden, wenn F keine Dichte besitzt, vgl. [Wol88, S.307]; das gilt auch für $E(X_E^2)$ und $V(X_E)$, die noch berechnet werden.)

Analog zu $E(X_E)$ kann $E(X_E^2)$ berechnet werden:

$$\begin{aligned} E(X_E^2) &= \int_0^\infty x^2 dF_E(x) = \\ &= \int_0^d x^2 f(x) dx + \underbrace{\int_d^\infty d^2 \cdot f(x) dx}_{=d^2 \cdot (1-F(d))} = \\ &= \int_0^d x^2 f(x) dx + d^2 \cdot (1 - F(d)). \end{aligned}$$

Auch $E(X_E^2)$ kann durch die Verteilungsfunktion F angegeben werden, wenn wiederum partiell integriert wird mit $u_1(x) = x^2$ und $v_1'(x) = f(x)$. Es ergibt sich dann:

$$\begin{aligned} E(X_E^2) &= \int_0^d x^2 f(x) dx + d^2 \cdot (1 - F(d)) = \\ &= \underbrace{x^2 F(x) \Big|_0^d}_{=d^2 \cdot F(d)} - \int_0^d 2x F(x) dx + d^2 \cdot (1 - F(d)) = \\ &= d^2 - 2 \int_0^d x F(x) dx. \end{aligned}$$

Für die Varianz von X_E ergibt sich aus den obigen Rechnungen:

$$\begin{aligned} V(X_E) &= E(X_E^2) - (E(X_E))^2 = \\ &= \int_0^d x^2 f(x) dx + d^2 \cdot (1 - F(d)) - \left(\int_0^d x f(x) dx + d \cdot (1 - F(d)) \right)^2 = \end{aligned}$$

$$\begin{aligned}
&= \int_0^d x^2 f(x) dx + d^2 \cdot (1 - F(d)) + \\
&\quad - \left(\left(\int_0^d x f(x) dx \right)^2 + 2d \cdot (1 - F(d)) \int_0^d x f(x) dx + (d \cdot (1 - F(d)))^2 \right) = \\
&= \int_0^d x^2 f(x) dx - \left(\int_0^d x f(x) dx \right)^2 + (1 - F(d)) \left(d^2 \cdot F(d) + 2d \cdot \int_0^d x f(x) dx \right).
\end{aligned}$$

Drückt man $E(X_E^2)$ und $E(X_E)$ nur durch die Verteilungsfunktion F aus, dann kann die Varianz von X_E folgendermaßen angegeben werden:

$$\begin{aligned}
V(X_E) &= E(X_E^2) - (E(X_E))^2 = \\
&= d^2 - 2 \int_0^d x F(x) dx - \left(d - \int_0^d F(x) dx \right)^2 = \\
&= d^2 - 2 \int_0^d x F(x) dx - d^2 + 2d \cdot \int_0^d F(x) dx - \left(\int_0^d F(x) dx \right)^2 = \\
&= 2 \int_0^d (d - x) F(x) dx - \left(\int_0^d F(x) dx \right)^2.
\end{aligned}$$

Die Schadenshöhen des Rückversicherers

Von einem Schaden X übernimmt der Rückversicherer den Anteil X_R , der durch $X_R = \max\{0, X - d\}$ angegeben werden kann. An Schäden, deren Schadenshöhe die Priorität d nicht übersteigt, ist der Rückversicherer nicht beteiligt, deshalb sind diese Schäden für den Rückversicherer nicht relevant und sie können bei der Erstellung eines Modells für den Rückversicherer nicht als Schäden betrachtet werden. Es werden also aus Sicht des Rückversicherers nur Schäden betrachtet, für die $X_R > 0$ gilt, das heißt, Schäden die die Bedingung $X > d$ erfüllen. Diese Schäden werden im folgenden mit X_d bezeichnet. X_d gibt dabei die Höhe des Schadens aus Sicht des Rückversicherers an, das heißt es gilt $X_d = X - d$. Es wird angenommen, dass d so gewählt wurde, dass die Wahrscheinlichkeit für den Eintritt eines Schadens, dessen Höhe die Priorität übersteigt, größer als 0 ist (d. h. $P(X > d) > 0$).

Für die Verteilungsfunktion F_d von X_d gilt (vgl. [Mac02, S. 336]):

$$F_d(x) = P(X_d < x) = P(X < x + d | X > d).$$

$P(X < x + d | X > d)$ lässt sich folgendermaßen berechnen:

$$P(X < x + d | X > d) = \frac{P(d < X < x + d)}{P(X > d)} = \frac{P(X < x + d) - P(X < d)}{P(X > d)}.$$

Deshalb gilt für F_d (F bezeichnet hier wiederum die Verteilungsfunktion von X):

$$F_d(x) = \frac{P(X < x + d) - P(X < d)}{P(X > d)} = \frac{F(x + d) - F(d)}{1 - F(d)}.$$

Besitzt F eine Dichte f , dann besitzt auch F_d eine Dichte f_d und es gilt:

$$\begin{aligned} f_d(x) &= F_d'(x) = \\ &= \left(\frac{F(x + d) - F(d)}{1 - F(d)} \right)' = \\ &= \frac{f(x + d)(1 - F(d))}{(1 - F(d))^2} = \\ &= \frac{f(x + d)}{1 - F(d)}. \end{aligned}$$

Für den Erwartungswert von X_d gilt, falls F eine Dichte f besitzt:

$$\begin{aligned} E(X_d) &= \int_0^\infty x f_d(x) dx = \\ &= \int_0^\infty x \frac{f(x + d)}{1 - F(d)} dx = \\ &= \frac{1}{1 - F(d)} \int_d^\infty (x - d) f(x) dx. \end{aligned}$$

Analog gilt für $E(X_d^2)$:

$$\begin{aligned} E(X_d^2) &= \int_0^\infty x^2 f_d(x) dx = \\ &= \int_0^\infty x^2 \frac{f(x + d)}{1 - F(d)} dx = \\ &= \frac{1}{1 - F(d)} \int_d^\infty (x - d)^2 f(x) dx = \\ &= \frac{1}{1 - F(d)} \int_d^\infty (x^2 - 2d \cdot x + d^2) f(x) dx. \end{aligned}$$

Daraus ergibt sich für die Varianz von X_d :

$$\begin{aligned}
V(X_d) &= E(X_d^2) - (E(X_d))^2 = \\
&= \frac{1}{1-F(d)} \int_d^\infty (x^2 - 2d \cdot x + d^2) f(x) dx - \left(\frac{1}{1-F(d)} \int_d^\infty (x-d) f(x) dx \right)^2 = \\
&= \frac{1}{1-F(d)} \left(\int_d^\infty x^2 f(x) dx - 2d \cdot \int_d^\infty x f(x) dx + d^2 \underbrace{\int_d^\infty f(x) dx}_{=1-F(d)} \right) + \\
&\quad - \left(\frac{1}{1-F(d)} \int_d^\infty x f(x) dx - \frac{1}{1-F(d)} d \cdot \underbrace{\int_d^\infty f(x) dx}_{=1-F(d)} \right)^2 = \\
&= \frac{1}{1-F(d)} \int_d^\infty x^2 f(x) dx - \frac{2d}{1-F(d)} \int_d^\infty x f(x) dx + d^2 + \\
&\quad - \left(\frac{1}{1-F(d)} \int_d^\infty x f(x) dx \right)^2 + \frac{2d}{1-F(d)} \cdot \int_d^\infty x f(x) dx - d^2 = \\
&= \frac{1}{1-F(d)} \int_d^\infty x^2 f(x) dx - \frac{1}{(1-F(d))^2} \left(\int_d^\infty x f(x) dx \right)^2.
\end{aligned}$$

Besitzt F keine Dichte bzw. ist die Dichte von F nicht bekannt, dann können der Erwartungswert und die Varianz von X_d auch ohne die Dichte f angegeben werden:

$$\begin{aligned}
E(X_d) &= \frac{1}{1-F(d)} \int_d^\infty (x-d) dF(x), \\
V(X_d) &= \frac{1}{1-F(d)} \int_d^\infty x^2 dF(x) - \frac{1}{(1-F(d))^2} \left(\int_d^\infty x dF(x) \right)^2.
\end{aligned}$$

Der Gesamtschaden von Erst- und Rückversicherer

Nachdem die Verteilungen der Schadenzahl und der Schadenshöhen für den Erstversicherer und für den Rückversicherer beschrieben wurden, sollen nun die Gesamtschäden von Erst- und Rückversicherer dargestellt werden. Am Beginn des Abschnitts wurde der Gesamtschaden des Erstversicherers angegeben mit:

$$S_E = \sum_{i=1}^N \min\{X_i, d\}.$$

Bezeichnet man mit $X_{i,E}$ jenen Anteil der Schadenshöhe X_i , den der Erstversicherer übernimmt (also $X_{i,E} = \min\{X_i, d\}$), dann lässt sich der Gesamtschaden des Erstver-

sicherers folgendermaßen angeben:

$$S_E = \sum_{i=1}^N X_{i,E}.$$

Mit den Formeln des kollektiven Modells ergibt sich daraus für den Erwartungswert und die Varianz von S_E :

$$\begin{aligned} E(S_E) &= E(N)E(X_E), \\ V(S_E) &= E(N)V(X_E) + V(N)(E(X_E))^2. \end{aligned}$$

Für den Gesamtschaden S_R des Rückversicherers gilt:

$$S_R = \sum_{i=1}^N \max\{X_i - d, 0\}.$$

Die Anzahl der Schäden, an denen der Rückversicherer beteiligt ist, lässt sich mit der eingeführten Zufallsvariable N_R angeben. Der Anteil $X_{i,R}$, den der Rückversicherer von der Schadenshöhe X_i übernimmt, lässt sich durch $X_{i,R} = \max\{X_i - d, 0\}$ ausdrücken. Hierbei werden aber auch Schäden erfasst, an denen der Rückversicherer nicht beteiligt ist, also Schäden mit $X_i \leq d$. Da Schäden, an denen der Rückversicherer nicht beteiligt ist, keine Schäden im kollektiven Modell des Rückversicherers darstellen sollen deshalb nur jene $X_{i,R}$ betrachtet werden, die $X_{i,R} > 0$ erfüllen (das sind jene X_i , die $X_i > d$ erfüllen). Diese werden mit $X_{1,d}, X_{2,d}, \dots$ durchnummeriert, der Gesamtschaden des Rückversicherers lässt sich dann folgendermaßen ausdrücken:

$$S_R = \sum_{i=1}^{N_R} X_{i,d}.$$

Für den Erwartungswert und die Varianz von S_R gilt:

$$\begin{aligned} E(S_R) &= E(N_R)E(X_d), \\ V(S_R) &= E(N_R)V(X_d) + V(N_R)(E(X_d))^2. \end{aligned}$$

Bei der hier beschriebenen Art der Modellierung wird der Gesamtschaden des Rückversicherers mit einem kollektiven Modell mit der Schadenzahl N_R und den Schadenshöhen $X_{1,d}, X_{2,d}, \dots$ modelliert. Eine andere Möglichkeit wäre gewesen, den Gesamtschaden des

Rückversicherers mit der Schadenzahl N und den Schadenshöhen $X_{1,R}, X_{2,R}, \dots$ (wobei $X_{i,R} = \max\{X_i - d, 0\}$) zu beschreiben, allerdings würden dann auch Schäden mit der Schadenshöhe 0 eingehen. Dies könnte beispielsweise zu Problemen führen, wenn im Rückversicherungsvertrag Fixkosten für die Abwicklung jedes einzelnen Schadens festgelegt wurden und das Modell zur Kalkulation von Zahlungen an den Rückversicherer verwendet werden soll (vgl. [Hei87, S. 162]).

Kumulschadenexzedentenrückversicherung

Diese Art der Rückversicherung funktioniert ähnlich wie die Schadenexzedentenrückversicherung, allerdings bezieht sich der Rückversicherungsschutz nicht auf einzelne Schäden, sondern auf Schadenereignisse. Ein Schadenereignis kann mehrere Verträge des Bestandes betreffen, beispielsweise kann das ein Hagelschauer oder Hochwasser sein. Die einzelnen Schadenereignisse werden wie Einzelschäden im kollektiven Modell erfasst, bezeichnet man mit der Zufallsvariable N^* die Anzahl der Schadenereignisse und mit $Y_1, Y_2, \dots$ die Schadenshöhen der einzelnen Schadenereignisse, dann ergibt sich der Gesamtschaden folgendermaßen:

$$S = \sum_{i=1}^{N^*} Y_i.$$

Zwischen dem Erstversicherer und dem Rückversicherer wird wie bei der Schadenexzedentenrückversicherung eine Priorität d^* festgelegt, bis zu der der Erstversicherer die Schäden eines Schadenereignisses vollständig begleicht, übersteigt die Schadenshöhe des Schadenereignisses die Priorität, dann übernimmt der Rückversicherer jenen Betrag, der die Priorität übersteigt (vgl. [Lie09, S. 175]). Von einem einzelnen Schadenereignis Y übernimmt der Erstversicherer also folgenden Betrag:

$$Y_E = \min\{Y, d^*\}.$$

Der Rückversicherer übernimmt den Anteil Y_R :

$$Y_R = \max\{Y - d^*, 0\}.$$

Für die jährlichen Gesamtschäden S_E und S_R von Erst- und Rückversicherer heißt das:

$$S_E = \sum_{i=1}^{N^*} \min\{Y_i, d\},$$

$$S_R = \sum_{i=1}^{N^*} \max\{Y_i - d, 0\}.$$

Die Modellierung kann analog zur Schadenexzedentenrückversicherung erfolgen.

Jahresüberschadenexzedentenrückversicherung

In der Jahresüberschadenexzedentenrückversicherung (auch als Stop-Loss-Rückversicherung bezeichnet) wird eine Priorität $h \in \mathbb{R}, h > 0$ festgelegt und der Jahresgesamtschaden S wird in folgender Weise zwischen Erst- und Rückversicherer aufgeteilt (vgl. [Lie09, S. 324]): Der Erstversicherer übernimmt den gesamten Betrag bis zur Gesamtschadenshöhe h :

$$S_E = \min\{S, h\}.$$

Der Rückversicherer übernimmt den Anteil des Gesamtschadens, der über die Priorität h hinaus geht:

$$S_R = \max\{S - h, 0\}.$$

Die Aufteilung des Gesamtschadens erfolgt also wie die Aufteilung von Einzelschäden in der Schadenexzedentenrückversicherung, deshalb können die dort getätigten Überlegungen zur Beschreibung des Jahresüberschadenexzedentenvertrages benutzt werden:

Ist F_S die Verteilung des Jahresgesamtschadens S , dann ist der Anteil S_E des Erstversicherers folgendermaßen verteilt:

$$F_{S_E}(x) = \begin{cases} F_S(x) & x < h, \\ 1 & x \geq h. \end{cases}$$

Für den Erwartungswert und die Varianz von S_E gilt:

$$\begin{aligned} E(S_E) &= h - \int_0^h F_S(x) dx, \\ V(S_E) &= 2 \int_0^h (h - x) F_S(x) dx - \left(\int_0^h F_S(x) dx \right)^2. \end{aligned}$$

Im Gegensatz zu den Einzelschäden muss der Gesamtschaden für den Rückversicherer auch betrachtet werden, wenn dieser kleiner als die Priorität ist. Die Wahrscheinlichkeit, dass der Gesamtschaden S nicht höher als die Priorität ist, ist durch $P(S \leq h) = F_S(h)$ gegeben. Die Verteilung F_{S_R} von S_R kann deshalb mit der Verteilung des Gesamtschadens

angegeben werden und es gilt (vgl. [Wol88, S. 306]):

$$F_{S_R}(x) = \begin{cases} F_S(x+h) & x \geq 0, \\ 0 & x < 0. \end{cases}$$

Für den Erwartungswert von S_R ergibt sich daraus:

$$\begin{aligned} E(S_R) &= \int_0^\infty x dF_S(x+h) = \\ &= \int_h^\infty (x-h) dF_S(x). \end{aligned}$$

Für $E(S_R^2)$ gilt:

$$\begin{aligned} E(S_R^2) &= \int_0^\infty x^2 dF_S(x+h) = \\ &= \int_h^\infty (x-h)^2 dF_S(x). \end{aligned}$$

Deshalb gilt für die Varianz von S_R :

$$\begin{aligned} V(S_R) &= E(S_R^2) - (E(S_R))^2 = \\ &= \int_h^\infty (x-h)^2 dF_S(x) - \left(\int_h^\infty (x-h) dF_S(x) \right)^2. \end{aligned}$$

Da S durch $S_E + S_R$ ausgedrückt werden kann gilt für $E(S)$:

$$E(S) = E(S_E) + E(S_R).$$

Der Erwartungswert von S_R kann deshalb auch folgendermaßen angegeben werden:

$$E(S_R) = E(S) - E(S_E) = E(S) - h + \int_0^h F_S(x) dx.$$

Diese Art der Rückversicherung kann sowohl im individuellen, als auch im kollektiven Modell angewendet werden.

5.1.3 Einsatzbereiche der verschiedenen Rückversicherungsarten

Für den Einsatz einer bestimmten Rückversicherungsart können unterschiedliche Gründe ausschlaggebend sein, dadurch ergeben sich verschiedene Versicherungssparten, in denen der Einsatz sinnvoll ist (nach [Sch98, S. 22]):

Die Quotenrückversicherung und die Jahresüberschadenexzedentenrückversicherung bieten Schutz vor Schwankungen bei kleinen und mittleren Schadenshöhen, Beispiele für den Einsatz sind die allgemeine Haftpflichtversicherung, die KFZ-Haftpflichtversicherung und die KFZ-Kaskoversicherung.

Im Gegensatz dazu bieten die Summenexzedentenrückversicherung und die Schadenexzedentenrückversicherung Schutz vor Großschäden, außerdem kann durch den Umstand, dass die Schadenshöhen für den Erstversicherer begrenzt werden, eine Homogenisierung der Schadenshöhen erreicht werden. Eingesetzt werden diese beiden Rückversicherungsarten beispielsweise in der Feuerversicherung und in der Industrieversicherung.

Die Kumulschadenexzedentenrückversicherung eignet sich für die Rückversicherung von Schadenereignissen, deren Gesamtschaden sich aus vielen, mitunter nur kleinen oder mittleren Schadenshöhen zusammensetzt, und dadurch dennoch sehr hoch sein kann. Praktisch eingesetzt wird die Kumulschadenexzedentenrückversicherung in der Hagelrückversicherung und in der Rückversicherung von Hochwasserschäden.

5.2 Selbstbehalte

Neben der Rückversicherung spielen Selbstbehalte (auch Franchisen genannt) eine große Rolle bei der Risikominderung. Die Selbstbehalte werden von den VersicherungsnehmerInnen getragen, im Gegensatz zur Rückversicherung schließt das Versicherungsunternehmen hier keine weiteren Verträge ab. Das Hauptziel, das mit Selbstbehalten erreicht werden soll, ist analog zu dem der Rückversicherung, nämlich die Teilung und Reduktion des versicherungstechnischen Risikos. Darüber hinaus kann mit manchen Selbstbehalten erreicht werden, dass Schäden mit kleiner Schadenshöhe vollständig von den betroffenen VersicherungsnehmerInnen getragen werden, wodurch für das Versicherungsunternehmen keine Kosten für die Bearbeitung dieser sogenannten „Bagatellschäden“ entstehen (vgl. [Hei87, S. 154]). Außerdem können VersicherungsnehmerInnen durch Selbstbehalte dazu angehalten werden, Schadenshöhen zu senken oder Schäden vollständig zu vermeiden (vgl. [DMS02, S. 128]).

In diesem Abschnitt werden die wichtigsten Formen der Selbstbeteiligung behandelt, es handelt sich dabei ausschließlich um sogenannte Kapitalfranchisen, bei denen Schäden

zwischen dem Versicherungsunternehmen und VersicherungsnehmerInnen aufgeteilt werden (eine andere Form wäre die Zeitfranchise, die aber nur dann angewendet wird, wenn Schäden über einen längeren Zeitraum verteilt anfallen, vgl. [DMS02, S. 128]).

Bei der folgenden Beschreibung der Franchisen bezeichnet die Zufallsvariable X stets die Schadenshöhe, die zwischen dem Versicherungsunternehmen und VersicherungsnehmerInnen aufgeteilt werden soll (dies kann sowohl die Höhe eines einzelnen Schadens, als auch der jährliche Gesamtschaden eines Versicherungsvertrages sein, je nachdem, worauf sich ein Selbstbehalt bezieht). Der Anteil des Versicherungsunternehmens wird mit X_U bezeichnet, der Anteil von VersicherungsnehmerInnen mit X_N . Die Darstellung der verschiedenen Franchisen erfolgt nach [Hei87, S. 154f].

Prozentuale Selbstbeteiligung

Bei der prozentualen Selbstbeteiligung erfolgt die Aufteilung zwischen VersicherungsnehmerIn und Versicherungsunternehmen folgendermaßen (mit $q \in (0, 1)$):

$$\begin{aligned}X_N &= q \cdot X, \\X_U &= (1 - q)X.\end{aligned}$$

Die prozentuale Selbstbeteiligung lässt sich mit der Quotenrückversicherung vergleichen: VersicherungsnehmerInnen sind hier unabhängig von der Höhe des Schadens mit einem festen Anteil an diesem beteiligt. Mit einer prozentualen Beteiligung sollen vor allem eine Senkung der Schadenshöhen bzw. die Vermeidung von Schäden erreicht werden (vgl. [Wag00, S. 317]).

Abzugsfranchise

Sei $a \in \mathbb{R}$, $a > 0$, dann lässt sich die Aufteilung des Schadens X nach der Abzugsfranchise folgendermaßen beschreiben:

$$\begin{aligned}X_N &= \min\{X, a\}, \\X_U &= \max\{X - a, 0\}.\end{aligned}$$

Die Berechnung des Selbstbehalts bei der Abzugsfranchise funktioniert nach dem selben Schema wie die Aufteilung der Schäden in der Schadenexzedentenrückversicherung: Bis zur Schadenshöhe a wird der Schaden vollständig von der Versicherungsnehmerin oder dem Versicherungsnehmer beglichen, alles was über die Schadenshöhe a hinausgeht wird

vom Versicherungsunternehmen übernommen. Praktisch wird die Abzugsfranchise in vielen Versicherungssparten eingesetzt, sie kann dazu genutzt werden, die Bearbeitung von kleinen Schäden durch das Versicherungsunternehmen zu vermeiden, dazu muss a so gewählt werden, dass es dem größtmöglichen „Bagatellschaden“ entspricht (vgl. [Sch94, S. 363]).

Integralfranchise

Bei der Integralfranchise wird der Schaden folgendermaßen aufgeteilt (dabei gilt: $a \in \mathbb{R}$, $a > 0$):

$$X_N = \begin{cases} X & X < a, \\ 0 & X \geq a. \end{cases}$$

$$X_U = \begin{cases} 0 & X < a, \\ X & X \geq a. \end{cases}$$

Schäden, deren Schadenshöhe kleiner als der Betrag a ist, müssen hier vollständig von der/dem VersicherungsnehmerIn beglichen werden. Alle Schäden, deren Schadenshöhe zumindest a beträgt werden hingegen vollständig vom Versicherungsunternehmen getragen. Die Integralfranchise wird nur selten angewendet, das Haupteinsatzgebiet liegt in der Transportversicherung, gebräuchlich ist sie hier vor allem beim Seetransport (vgl. [FW01, S. 257]).

Verswindende Abzugsfranchise

Seien $a, b \in \mathbb{R}$ mit $0 < a < b$. Dann kann die Aufteilung des Schadens X mit der verschwindenden Abzugsfranchise folgendermaßen angegeben werden:

$$X_N = \begin{cases} X & x \leq a, \\ a \frac{b-X}{b-a} & a < X \leq b, \\ 0 & X > b. \end{cases}$$

$$X_U = \begin{cases} 0 & x \leq a, \\ \frac{b}{b-a}(X - a) & a < X \leq b, \\ X & X > b. \end{cases}$$

Die verschwindende Abzugsfranchise kann als Kombination aus der Abzugsfranchise und der Integralfranchise gesehen werden (vgl. [DMS02, S. 140]): Bis zur Schadenshöhe a wird der Schaden vollständig von der/dem VersicherungsnehmerIn getragen. Ist die Schadenshöhe zwischen dem Betrag a und dem Betrag b , dann übernimmt das Versicherungsunternehmen den Anteil $\frac{b}{b-a}(X - a)$, VersicherungsnehmerInnen übernehmen den Anteil $a\frac{b-X}{b-a}$. Da $b > a$ gilt nimmt für VersicherungsnehmerInnen mit steigender Schadenshöhe der übernommene Anteil ab, für das Versicherungsunternehmen hingegen nimmt er zu. Ab der Schadenshöhe a übernimmt das Versicherungsunternehmen den Schaden vollständig. Eingesetzt wird die verschwindende Abzugsfranchise vor allem in der industriellen Feuerversicherung (vgl. [Sch94, S. 365]).

6 Zusammenfassung und Ausblick

In dieser Arbeit zur Schadenversicherungsmathematik wurden neben versicherungstechnischen und mathematischen Grundlagen die Modellierung von Risiken, die Berechnung von Versicherungsprämien und die Minderung des versicherungstechnischen Risikos behandelt. Dies stellt allerdings nur einen Teil der Schadenversicherungsmathematik dar. Deshalb sollen zum Schluss noch einige weitere Anwendungen der Mathematik in der Schadenversicherung genannt werden, auf die hier nicht eingegangen wurde (diese Liste kann selbstverständlich nur als Auswahl betrachtet werden):

- Risikoprozesse: In dieser Arbeit wurden Größen wie die Schadenszahl oder der Gesamtschaden stets auf eine gesamte Periode bezogen, praktisch ist es aber auch von Bedeutung, wie viele Schäden nach einer gewissen Zeit t (z. B. nach einem Monat) aufgetreten sind und wie hoch der Gesamtschaden dann ist. Dies wird mit Risikoprozessen wie dem Gesamtschadenprozess oder dem Schadenszahlprozess beschrieben (vgl. [13, S. 36]).
- Ruintheorie: Die Ruintheorie untersucht die Wahrscheinlichkeit des sogenannten technischen Ruins (dieser tritt dann ein, wenn die Höhe der aufgetretenen Schäden die Summe aus Prämieinnahmen und einem eventuell vorhandenen Sicherheitskapital übersteigt) und versucht diese zu minimieren (vgl. [13, S. 39]).
- Credibility-Theorie: Die Credibility-Theorie beschäftigt sich mit der Festlegung der Prämie als gewichtetes Mittel aus der individuellen Schadenerfahrung eines Risikos und dem Schadenbedarf des gesamten Bestandes (vgl. [Hei87, S. 138]).
- Spätschadenreservierung: Nicht alle aufgetretenen Schäden werden sofort bemerkt oder dem Versicherungsunternehmen gemeldet. Die Schadenreservierung beschäftigt sich mit der Reservierung von Kapital für die Regulierung dieser Schäden (vgl. [Sch02, S. 269]).

Literaturverzeichnis

- [AS88] ALBRECHT, Peter; SCHWAKE, Edmund: Das versicherungstechnische Risiko. In: FARNY, Dieter (Hrsg.); HELTEN, Elmar (Hrsg.); KOCH, Peter (Hrsg.) ; SCHMIDT, Reimer (Hrsg.): *Handwörterbuch der Versicherung*. Karlsruhe: Verlag Versicherungswirtschaft, 1988, S. 650–657.
- [Asm03] ASMUSSEN, Søren: *Applied probability and queues*. New York: Springer-Verlag, 2003.
- [AW05] ADELMEYER, Moritz; WARMUTH, Elke: *Finanzmathematik für Einsteiger: Von Anleihen über Aktien zu Optionen*. Wiesbaden: Vieweg, 2005.
- [BF87] BERTRAM, Jürgen; FEILMEIER, Manfred: *Anwendung numerischer Methoden in der Risikothorie*. Karlsruhe: Verlag Versicherungswirtschaft, 1987.
- [BGH⁺97] BOWERS, Newton L.; GERBER, Hans U.; HICKMAN, James C.; JONES, Donald A. ; NESBITT, Cecil J.: *Actuarial Mathematics*. Schaumburg: Society of Actuaries, 1997.
- [Büh88] BÜHLMANN, Hans: Entwicklungstendenzen in der Risikothorie. In: *Jahresbericht der Deutschen Mathematiker-Vereinigung 90* (1988), Nr. 3, S. 111–128.
- [DMS02] DIENST, Hans R.; MACK, Thomas ; SEGERER, Günther: Risikoteilung. In: HELBIG, Manfred (Hrsg.): *Beiträge zum versicherungsmathematischen Grundwissen*. Karlsruhe: Verlag Versicherungswirtschaft, 2002, S. 121–148.
- [Far06] FARNY, Dieter: *Versicherungsbetriebslehre*. Karlsruhe: Verlag Versicherungswirtschaft, 2006.
- [FW01] FÜRSTENWERTH, Jörg Frank v.; WEISS, Alfred: *Versicherungsalphabet: Begriffserläuterungen der Versicherung aus Theorie und Praxis*. Karlsruhe: Verlag Versicherungswirtschaft, 2001.

- [Gru96] GRUNDMANN, Wolfgang: *Finanz- und Versicherungsmathematik*. Stuttgart; Leipzig: Teubner, 1996.
- [Haf89] HAFNER, Robert: *Wahrscheinlichkeitsrechnung und Statistik*. Wien; New York: Springer-Verlag, 1989.
- [Hei87] HEILMANN, Wolf-Rüdiger: *Gundbegriffe der Risikotheorie*. Karlsruhe: Verlag Versicherungswirtschaft, 1987.
- [Hei88] HEILMANN, Wolf-Rüdiger: Schadenversicherungsmathematik. In: FARNY, Dieter (Hrsg.); HELTEN, Elmar (Hrsg.); KOCH, Peter (Hrsg.) ; SCHMIDT, Reimer (Hrsg.): *Handwörterbuch der Versicherung*. Karlsruhe: Verlag Versicherungswirtschaft, 1988, S. 753–758.
- [Hip90] HIPPE, Christian: *Risikotheorie: Stochastische Modelle und statistische Methoden*. Karlsruhe: Verlag Versicherungswirtschaft, 1990.
- [HK91] HELTEN, Elmar; KARTEN, Walter: Das Risiko und seine Kalkulation. In: GROSSE, Walter (Hrsg.): *Versicherungszyklopädie Band 2*. Wiesbaden: Gabler, 1991, S. 223–275.
- [Koc88] KOCH, Peter: Geschichte der Versicherung. In: FARNY, Dieter (Hrsg.); HELTEN, Elmar (Hrsg.); KOCH, Peter (Hrsg.) ; SCHMIDT, Reimer (Hrsg.): *Handwörterbuch der Versicherung*. Karlsruhe: Verlag Versicherungswirtschaft, 1988, S. 223–232.
- [Koc05] KOCH, Peter: *Versicherungswirtschaft: Ein einführender Überblick*. Karlsruhe: Verlag Versicherungswirtschaft, 2005.
- [Kre85] KREMER, Erhard: *Einführung in die Versicherungsmathematik*. Göttingen: Vandenhoeck und Ruprecht, 1985.
- [KS88] KUON, Sigfried; STICKER, Klaus: Bonus-/Malus-System. In: FARNY, Dieter (Hrsg.); HELTEN, Elmar (Hrsg.); KOCH, Peter (Hrsg.) ; SCHMIDT, Reimer (Hrsg.): *Handwörterbuch der Versicherung*. Karlsruhe: Verlag Versicherungswirtschaft, 1988, S. 91–94.
- [Lie09] LIEBWEIN, Peter: *Klassische und moderne Formen der Rückversicherung*. Karlsruhe: Verlag Versicherungswirtschaft, 2009.

- [Mac02] MACK, Thomas: *Schadenversicherungsmathematik*. Karlsruhe: Verlag Versicherungswirtschaft, 2002.
- [Mik09] MIKOSCH, Thomas: *Non-Life Insurance Mathematics. An introduction with the Poisson Process*. Berlin; Heidelberg: Springer-Verlag, 2009.
- [Pan81] PANJER, Harry H.: Recursive evaluation of a family of compound distributions. In: *ASTIN Bulletin* (1981), Nr. 12, S. 22–26.
- [PW92] PANJER, Harry H.; WILLMOT, Gordon E.: *Insurance risk models*. Schaumburg: Society of Actuaries, 1992.
- [Sch94] SCHRADIN, Heinrich: *Erfolgsorientiertes Versicherungsmanagement: Betriebswirtschaftliche Steuerungskonzepte auf risikotheorietischer Grundlage*. Karlsruhe: Verlag Versicherungswirtschaft, 1994.
- [Sch95] SCHRÖTER, Klaus J.: *Verfahren zur Approximation der Gesamtschadenverteilung: Systematisierung, Techniken und Vergleiche*. Karlsruhe: Verlag Versicherungswirtschaft, 1995.
- [Sch97] SCHOTT, Winfried: *Preise für versicherungstechnische Risiken: Zu Rationalität und Realitätsnähe bei Anwendung aktuarieller und ökonomischer Modelle*. Karlsruhe: Verlag Versicherungswirtschaft, 1997.
- [Sch98] SCHRADIN, Heinrich: *Finanzielle Steuerung der Rückversicherung unter besonderer Berücksichtigung von Großschadenereignissen und Fremdwährungsrisiken*. Karlsruhe: Verlag Versicherungswirtschaft, 1998.
- [Sch02] SCHMIDT, Klaus D.: *Versicherungsmathematik*. Berlin; Heidelberg: Springer-Verlag, 2002.
- [Sch05] SCHULENBURG, Johann-Matthias von d.: *Versicherungsökonomik : Ein Leitfaden für Studium und Praxis*. Karlsruhe: Verlag Versicherungswirtschaft, 2005.
- [Str88a] STRAUSS, Jürgen: Die Rückversicherungsverfahren in der Praxis. In: DIENST, Hans-Rudolf (Hrsg.): *Mathematische Verfahren der Rückversicherung*. Karlsruhe: Verlag Versicherungswirtschaft, 1988, S. 7–34.
- [Str88b] STRAUB, Erwin: *Non-Life Insurance Mathematics*. Berlin: Springer-Verlag, 1988.

- [Sun99] SUNDT, Bjørn: *An Introduction to Non-Life Insurance Mathematics*. Karlsruhe: Verlag Versicherungswirtschaft, 1999.
- [Wag00] WAGNER, Fred: *Risk Management im Erstversicherungsunternehmen: Modelle, Strategien, Ziele, Mittel*. Karlsruhe: Verlag Versicherungswirtschaft, 2000.
- [Wol88] WOLFSDORF, Kurt: *Versicherungsmathematik, Teil 2. Theoretische Grundlagen, Risikotheorie, Sachversicherung*. Stuttgart: Teubner, 1988.

Internetquellen

- [1] *Inverse Normalverteilung.* URL: http://de.wikipedia.org/wiki/Inverse_Normalverteilung, zuletzt zugegriffen am: 20.11.2009.
- [2] *Jahresbericht der Finanzmarktaufsichtsbehörde 2008: Aufsicht über Versicherungsunternehmen.* URL: http://www.fma.gv.at/JBInteraktiv/2008/DE/CH0188_CMS1239137893647_detail.htm, zuletzt zugegriffen am: 16.01.2010.
- [3] *Versicherungsvertragsgesetz (VersVG) online.* URL: [http://www.jusline.at/Versicherungsvertragsgesetz_\(VersVG\).html](http://www.jusline.at/Versicherungsvertragsgesetz_(VersVG).html), zuletzt zugegriffen am: 20.01.2010.
- [4] BARTELS, Sven: *Nullnutzen-, Exponential- und Escherprinzip.* URL: <http://www.grin.com/e-book/99319/nullnutzen-exponential-und-escherprinzip>, zuletzt zugegriffen am: 05.09.2009.
- [5] DREES, Holger: *Risikothorie.* URL: [http://www.math.uni-hamburg.de/home/drees/VM2_05/VM II SS 05.pdf](http://www.math.uni-hamburg.de/home/drees/VM2_05/VM%20II%20SS%2005.pdf), zuletzt zugegriffen am: 09.07.2009.
- [6] GROLL, Andreas: *Modellierung und Bewertung von Katastrophenanleihen.* URL: <http://www.mathematik.uni-muenchen.de/~sekrfil/PDFs/Groll.pdf>, zuletzt zugegriffen am: 26.11.2009.
- [7] HIPPE, Christian: *Risikothorie 1.* URL: [http://insurance.fbv.uni-karlsruhe.de/rd_download/Risikothorie_1_\(Skript_2006_korrigiert\).pdf](http://insurance.fbv.uni-karlsruhe.de/rd_download/Risikothorie_1_(Skript_2006_korrigiert).pdf), zuletzt zugegriffen am: 29.09.2009.
- [8] HIPPE, Christian: *Spezialwissen Schadenversicherungsmathematik.* URL: http://insurance.fbv.uni-karlsruhe.de/rd_download/Skript_Schaden.pdf, zuletzt zugegriffen am: 29.09.2009.
- [9] KLÜPPELBERG, Claudia: *Prämienkalkulation.* URL: <http://www-m4.ma.tum.de/courses/SS04/riskth/Ri9.pdf>, zuletzt zugegriffen am: 05.09.2009.

- [10] MACK, Thomas: *Vorlesung „Schadenversicherungsmathematik“*. URL: <http://www-m4.ma.tum.de/courses/WS08-09/schaden/Vorlesung2008.pdf>, zuletzt zugegriffen am: 18.08.2009.
- [11] MINDLIN, Ilja: *Prämienberechnungsprinzipien*. URL: <http://www.mi.uni-koeln.de/~jeisenbe/Praemienberechnung.pdf>, zuletzt zugegriffen am: 26.08.2009.
- [12] PFEIFER, Dietmar: *Versicherungsmathematik I*. URL: <http://www.staff.uni-oldenburg.de/dietmar.pfeifer/VersMathI.pdf>, zuletzt zugegriffen am: 16.10.2009.
- [13] RIEDLE, Markus: *Script zu Risikotheorie*. URL: <http://www.math.hu-berlin.de/~riedle/research/scriptrisiko.pdf>, zuletzt zugegriffen am: 10.06.2009.
- [14] STANGL, Katharina: *Versicherungsmathematische Analyse des Bonus-Malus Systems in der KFZ-Haftpflichtversicherung*. URL: <http://www.stat.tugraz.at/dthesis/stangl07.pdf>, zuletzt zugegriffen am: 12.10.2009.

Abbildungsverzeichnis

2.1	Die Dichtefunktion der Normalverteilung mit $\mu = 0$ und verschiedenen Werten für σ^2	19
3.1	Die Dichtefunktion der Gammaverteilung mit $\beta = 1$ und verschiedenen Werten für α	34
3.2	Die Dichtefunktion der Exponentialverteilung mit verschiedenen Werten für λ	41
3.3	Die Dichtefunktion der Inversen Normalverteilung mit $\mu = 1$ und verschiedenen Werten für λ	43
3.4	Die Dichtefunktion der Paretoverteilung mit $a = 1$ und verschiedenen Werten für b	45
3.5	Die Dichtefunktion der Logarithmischen Normalverteilung mit $\mu = 0$ und verschiedenen Werten für σ^2	53

Kurzzusammenfassung

Das Versicherungswesen beruht auf der Idee, dass viele Personen Geld bezahlen, um im Falle eines bestimmten Ereignisses Zahlungen zu erhalten, mit denen entstandene Schäden bzw. Nachteile ausgeglichen werden können. Schäden und damit Ausgaben für das Versicherungsunternehmen aber treten zufällig auf, während die Einnahmen in Form von Versicherungsprämien im Voraus bestimmt werden müssen. Mit Hilfe der Mathematik soll versucht werden, das Risiko für das Versicherungsunternehmen berechnen- bzw. abschätzbar zu machen, damit die Einnahmen des Versicherungsunternehmens mit einer hohen Wahrscheinlichkeit zumindest so hoch sind wie die Ausgaben.

In dieser Arbeit werden grundlegende mathematische Prinzipien der Schadenversicherung dargestellt. Der Schwerpunkt wird dabei auf den Umgang des Versicherungsunternehmens mit den übernommenen Risiken gelegt, konkret werden die wahrscheinlichkeitstheoretischen Modelle der Schadenversicherung, die Kalkulation von Prämien und Möglichkeiten zur Minderung des übernommenen Risikos thematisiert: Am Beginn der Arbeit wird das Versicherungswesen mit seinen Grundlagen behandelt, es wird dabei speziell auf das versicherungstechnische Risiko und den Ausgleich im Kollektiv eingegangen. Die Modellbildung in der Schadenversicherung kann entweder nach dem individuellen Modell, das von den jährlichen Gesamtschäden der einzelnen Versicherungsverträge ausgeht, oder nach dem kollektiven Modell, bei dem die Schadenshöhen einzelner Schäden erfasst werden, erfolgen. Für beide Modelle werden geeignete Wahrscheinlichkeitsverteilungen eingeführt und es wird gezeigt wie die Verteilung des Gesamtschadens ermittelt werden kann. Zur Berechnung von Prämien werden Prämienprinzipien, bei denen die Prämien anhand bestimmter Eigenschaften eines Versicherungsvertrages festgelegt werden, und die Erfahrungstarifizierung, die den individuellen Schadenverlauf eines Versicherungsvertrages zur Berechnung der Prämie heranzieht, behandelt. Die Minderung des versicherungstechnischen Risikos wird durch die Rückversicherung, bei der Versicherungsunternehmen selbst Versicherungsverträge abschließen um sich gegen bestimmte Gefahren zu schützen, und durch die Selbstbeteiligung von VersicherungsnehmerInnen an aufgetretenen Schäden thematisiert.

Abstract

This thesis describes the mathematical principles of non-life insurance. In the description it focuses on the handling of risks by the insurance company.

First the basic principles of insurance are illustrated. Then insurance risk models are discussed. Risk models for insurances can be divided into individual risk models and collective risk models. Individual risk models use claim distributions of individual insurance contracts to calculate the total claim distribution, while the total claim distribution is estimated from the distributions of individual claim amounts and the claim number distribution in collective risk models. For both the individual risk models and the collective risk models it is shown how the total claim can be calculated. The calculation of insurance premiums is demonstrated by the use of premium principles and experience rating. Premium principles determine a premium from certain characteristics of an insurance contract, while experience rating includes the individual claim history of a contract into the premium calculation. Insurance companies try to minimize the so-called “underwriting risk” by transferring risks. Risk transfer includes reinsurance, where the insurers conclude reinsurance contracts with reinsurers, and deductibles (excesses), where policyholders cover parts of their own claims.

Lebenslauf

Persönliche Daten

Name: Manuel Feindert
Geburtsdatum: 28.02.1984
Geburtsort: Grieskirchen
Eltern: Helga und Johann Feindert
Staatsbürgerschaft: Österreich

Bildungsgang

1990-1994: Volksschule in Neukirchen am Walde
1994-1998: Hauptschule in Neukirchen am Walde
1998-2003: HTL Leonding (Ausbildungszweig: EDV und Organisation)
17.06.2003: Matura mit gutem Erfolg bestanden
WS 2004/05: Beginn des Lehramtsstudiums (UF Mathematik, UF Physik) an der Universität Wien

Berufserfahrung

- zwischen 1999 und 2002: mehrere Ferialpraktika bei der Firma Gerstl Bauunternehmung in Wels; Aufgabenbereich: Administration der EDV-Anlagen
- Juli und August 2003: Ferialpraktikum im GRZ IT Center Linz; Aufgabenbereich: Programmierung kaufmännischer Software
- Oktober 2003 bis September 2004: Ableistung des Zivildienstes bei der Caritas Oberösterreich in einer Betreuungseinrichtung für psychisch kranke Menschen
- Sommer 2006, Sommer 2007 und Sommer 2008: Arbeit als Campbetreuer und Camplehrer bei Sommercamps der Kinderwelt Niederösterreich