

Beiträge zur Risikoforschung

Fachliche Grundlagen der Arbeiten des Instituts für Vermögensaufbau

Herausgeber
Institut für Vermögensaufbau (IVA) AG
November 2012

Beiträge zur Risikoforschung

Fachliche Grundlagen der Arbeiten des Instituts für Vermögensaufbau

Herausgeber:

Institut für Vermögensaufbau (IVA) AG
Nymphenburger Straße 113
D-80636 München
Tel +49 (0)89 461 391 70
Fax +49 (0)89 461 391 79
www.institut-va.de
mail@institut-va.de

Textbeiträge:

Dr. Gabriel Layes, Andreas Ritter, Dr. Andreas Beck

Datenanalysen:

Andreas Ritter, Benjamin Kant

Inhaltsverzeichnis

1	Einführung	4
2	Fachliche Grundlagen der Risikoanalyse	7
2.1	Grundüberlegungen	7
2.2	Risikoarten	8
2.3	Einfache Risikokennzahlen	9
2.4	Verteilungsbasierte Risikomodellierung	12
2.4.1	Verteilungsarten	12
2.4.2	Simulation von Aktienmarktrenditen	14
2.4.3	Simulation von Zinszeitreihen	17
2.4.4	Berücksichtigung von Ausfallwahrscheinlichkeiten	19
3	Bewertungsprozesse	22
3.1	Klassifizierung von Wertpapieren anhand ihrer Risikomerkmale	22
3.2	Bewertung von Wertpapierfonds / Sammelvermögen	24
3.3	Bewertung von Anleihen	24
3.4	Bewertung von Einzelaktien	25
3.5	Bewertung von Derivaten	26
3.6	Exemplarische Bewertung unterschiedlicher Finanzinstrumente	27
4	Risikobewertung bei langer Haltedauer	29
4.1	Quantitative Merkmale mehrperiodiger Haltedauer	29
4.2	Berücksichtigung qualitativer zukünftiger Szenarien	30

Abbildungsverzeichnis

1	VaR und CVaR bei zwei unterschiedlich riskanten Finanzprodukten	12
2	Simuliert normalverteilte Renditen im Vergleich zu realen Renditen	13
3	Mischverteilungen aus zwei Normalverteilungen	15
4	GARCH-simulierte Renditen im Vergleich zu realen Renditen	16
5	Simulierte Pfade einer Zero Rate aus der Zinsstrukturkurve für Bundesanleihen	18
6	Simulierte Verteilung einer Zero Rate aus der Zinsstrukturkurve für Bundesanleihen nach einem Jahr	18
7	Regressionsmodell zur Schätzung von Ausfallwahrscheinlichkeiten bei fehlender CDS-Prämie	21
8	Übersicht zur Klassifizierung der einzelnen Wertpapierarten	23
9	Schematische Darstellung des Bewertungsprozesses für Anleihen	24
10	Schematische Darstellung des Bewertungsprozesses für Einzelaktien	25
11	Schematische Darstellung des Bewertungsprozesses für Derivate	26
12	Simulierte Wahrscheinlichkeitsverteilungen von Finanzinstrumenten mit unterschiedlichem Ausfallrisiko	28
13	Schwankungsrisiko verschiedener Finanzinstrumente in Abhängigkeit von der Haltedauer	30

1 Einführung

Im Bereich “Risikoforschung” beschäftigt sich das Institut für Vermögensaufbau unter den folgenden drei Aspekten mit risikobehafteten Finanzprodukten und deren Zusammenstellung zu Wertpapierportfolios:

Analyse von Anlagemöglichkeiten und Finanzdienstleistungen für Privatanleger,

Bewertung der Anlagemöglichkeiten und Finanzdienstleistungen unter dem Aspekt der Eignung zur Verfolgung unterschiedlicher Anlageziele,

Vermittlung der Ergebnisse auf verständliche Art in Form von Studien, Zertifikaten oder anderen Dokumenten.

Das heißt, es werden risikobehaftete Finanzprodukte unter dem Aspekt des Verhältnisses von Chancen und Risiken analysiert, kategorisiert und für Privatanleger transparent gemacht.

Wie sich in dieser Dreiteilung bereits andeutet, werden sinnvolle Aussagen über Finanzprodukte unseres Erachtens erst dann möglich, wenn die Analyse, Bewertung und Vermittlung von Produkteigenschaften als drei verschiedene Aufgaben verstanden werden. Dieser Hinweis erscheint uns deshalb als nicht trivial, weil diese drei Ebenen in der öffentlichen Diskussion über Finanzprodukte häufig vermischt werden. Beispielsweise dann, wenn der Versuch unternommen wird, Finanzprodukte “an sich” als gut oder schlecht zu klassifizieren. Dahinter verbirgt sich oftmals ein stark verkürztes Verständnis von “Risiko”, welches das Risiko eines Finanzprodukts mit dem Ausmaß seiner Preisschwankungen über kurze Zeiträume gleichsetzt. Die Bewertung eines Finanzprodukts ist letztlich nur in Bezug auf verschiedene Anlageziele und deren Fristigkeit möglich und auch nur dann, wenn neben dem Marktrisiko des Produkts sein gesamtes Risikoprofil bekannt ist. Dazu müssen auf der Ebene der Risikomessung zunächst verschiedene Risikoarten unterschieden werden.

Bezüglich der Rendite r_i , die bei Investition in ein risikobehaftetes Finanzprodukt erwartet werden kann, gehen wir im Sinne multifaktorieller Kapitalmarktmodelle (Chen et al. 1986) davon aus, dass deren systematisch bestimmbarer Teil aus zwei Komponenten besteht: Dem risikolosen Zins Z_0 und der Summe der Risikoprämien, die dafür erwartet werden können, dass man sich durch die Investition verschiedenen Risikofaktoren RF_i im Ausmaß β aussetzt:

$$r_i = Z_0 + \beta_1 RF_1 + \dots + \beta_n RF_n$$

Handelt es sich bei dem Finanzprodukt um ein aktiv gemanagtes Produkt, kann zu diesen beiden Komponenten eine dritte hinzukommen, das sogenannten “Alpha”. Dabei handelt es sich um eine von Marktfaktoren unabhängige Renditekomponente, die ausschließlich die Managementaktivitäten widerspiegelt. Die obige Formel lässt sich somit folgendermaßen erweitern:

$$r_i = Z_0 + \beta_1 RF_1 + \dots + \beta_n RF_n + \alpha$$

Auch das Alpha lässt sich als Risikoprämie interpretieren, und zwar dafür, dass man sich durch die Investition in ein aktiv gemanagtes Produkt dem Managementrisiko aussetzt. Entsprechend kann das Alpha ebenso wie die Marktrisikoprämien sowohl positiv als auch negativ werden. Es gibt allerdings zwei wesentliche Unterschiede zwischen dem Alpha und den Marktrisikoprämien: Während letztere langfristig einen positiven Erwartungswert besitzen, ist der Erwartungswert des Alphas über die Menge aller aktiv gemanagten Produkte eines Marktes gleich Null, da das positive Alpha eines Managers zwingend das negative Alpha eines anderen Managers sein muss. Zum anderen ist der “Einsatz”, der gezahlt werden muss, um eine Chance auf die Vereinnahmung der jeweiligen Risikoprämie zu haben, bei den Marktrisikoprämien deutlich niedriger als beim Alpha, für das als “Einsatz” in aller Regel eine vergleichsweise hohe Managementgebühr gezahlt werden muss.

Dies führt vor Augen, dass der Chance auf die Vereinnahmung von Risikoprämien immer verschiedene renditereduzierende Faktoren gegenüber stehen: Das sind zum einen die Kosten für den Erwerb, die

Verwahrung und ggf. das aktive Management des Produkts, sowie desweiteren Steuern auf Kapitalerträge und die Inflation, die bei zunehmender Haltedauer ein immer wichtigerer Faktor wird. Für einen langfristig orientierten Anleger, der seine tatsächlich erwirtschaftete, **reale Rendite** r_{real} bestimmen will, stellt sich die obige Gleichung somit so dar, dass er von der erzielten Rendite r_i die Summe aller Kosten, die Steuern, sowie die Inflation über die Haltedauer abziehen muss:

$$r_{real} = (Z_0 + \beta_1 RF_1 + \dots + \beta_n RF_n + \alpha) - \left(\sum \text{Kosten} + \text{Steuern} + \text{Inflation} \right)$$

Wie stellt sich diese Gleichung dar, wenn man sie unter den heutigen Kapitalmarktbedingungen mit konkreten Zahlen befüllt? - Betrachten wir dazu zunächst den einfachen Fall eines maximal risikoaversen Anlegers, der sich zur größtmöglichen Vermeidung von Risiken und Kosten ausschließlich auf die Vereinnahmung risikoloser Zinserträge beschränkt. Er setzt somit gewissermaßen in der obigen Gleichung β , α und die Kosten auf Null, so dass nur noch die nominalen Zinszahlungen Z_0 den Steuern und der Inflation gegenüber stehen. Trifft man nun die Annahme, dass es der Europäischen Zentralbank gelingt, ihr Inflationsziel für den Euroraum von knapp 2% langfristig zu realisieren, muss dieser Anleger vor Steuern regelmäßige Zinszahlungen in Höhe von 2,6% p.a. erhalten, um nach Steuern zumindest die Kaufkraft der angelegten Beträge zu erhalten¹. Liegt die langfristige Inflationsrate nur ein Prozent höher und somit bei 3%, so sind vor Steuern bereits regelmäßige Zinszahlungen in Höhe von 4,1% p.a. notwendig, um nur die Kaufkraft der angelegten Beträge zu erhalten.

Versteht man unter dem risikolosen Zins Z_0 die Rendite langlaufender Staatsanleihen höchster Bonität und vergleicht diese mit den oben genannten Zahlen, so wird schnell klar, dass beim völligen Verzicht auf die Inkaufnahme von Risiken mit hoher Wahrscheinlichkeit langfristig noch nicht einmal der reale Kapitalerhalt und somit das Bescheidenste aller Anlageziele erreicht wird. Unter den derzeit herrschenden Kapitalmarktbedingungen stellt sich somit selbst für einen maximal risikoaversen Anleger mit Blick auf langfristige Anlageziele weniger die Frage, **ob** er Risiken eingehen sollte, sondern **welche er wie stark** eingehen sollte. Ein völliger Verzicht auf das Eingehen von Risiken ist ohnehin gar nicht möglich, da selbst der völlige Verzicht auf das Eingehen von Marktrisiken nur bedeutet, dass man sich sehr stark gegenüber dem Inflationsrisiko exponiert.

Insofern ist es bedauerlich, dass in den letzten Jahren unter dem Eindruck der Banken- und Staatsschuldenkrise die ohnehin nur schwach ausgeprägte Wertpapierkultur in Deutschland noch weiter gelitten hat. In breiten Bevölkerungsschichten scheint sich die Meinung durchgesetzt zu haben, risikobehaftete Finanzinstrumente eigneten sich nur für "Zocker", nicht für verantwortungsbewusste Menschen mit "seriösen" Anlagezielen. An dieser Entwicklung sind die Banken durch ihre Art des Umgangs mit Risiken in der Vergangenheit sicherlich keineswegs unschuldig. Allerdings schüttet man das Kind mit dem Bade aus, wenn man deshalb (potenziellen) Anlegern gegenüber die Investition in risikobehaftete Finanzinstrumente grundsätzlich erschwert oder dämonisiert.

Wir sind stattdessen der Auffassung, dass der grundsätzliche Nutzen risikobehafteter Finanzinstrumente auch im Lichte der vergangenen Jahre fachlich unbestreitbar ist, so dass eine sachgerechte Diskussion über die Risiken von Finanzinstrumenten stattfinden muss, die nur dann möglich ist, wenn die oben eingeführte Unterscheidung zwischen Fragen a) der Risikoanalyse, b) der Risikobewertung und c) der Risikoauflklärung klar getroffen wird.

Auf der Ebene der Risikoauflklärung haben die vergangenen Jahre offen gelegt, dass viele Anleger in mangelhafter Art und Weise über die Risiken von Finanzprodukten aufgeklärt worden sind. Auf der Suche nach den Ursachen stößt man allerdings auch hier auf etwas komplexere Zusammenhänge als das, was sich in der öffentlichen Diskussion widerspiegelt: Ein Teil der Probleme lässt sich wohl tatsächlich damit erklären, dass Risikoauflklärung von einigen Bankberatern im Lichte von Vertriebszielen nur in unzureichender Weise betrieben worden ist. In den meisten Fällen mangelhafter Risikoauflklärung muss man allerdings vermuten, dass sich auch Bankberater des Ausmaßes bestimmter Risiken nicht bewusst

¹ Unterstellt wird eine steuerliche Belastung in Höhe von 26,4%, entsprechend der Abgeltungssteuer in Höhe von 25% plus 5,5%igem Solidaritätszuschlag. Ist der Anleger Kirchenmitglied, steigt diese Steuerlast je nach Bundesland noch weiter.

waren, so dass ihnen schwerlich vorzuwerfen ist, hierüber nicht genügend aufgeklärt zu haben. Dies gilt etwa für das Emittentenrisiko von Zertifikaten oder das Liquiditätsrisiko von Immobilienfonds, das offenbar auch von zahlreichen professionellen Marktteilnehmern unterschätzt worden ist und somit eher ein Problem der Risikobewertung darstellt.

Das Gleiche gilt für einige in den Risikoklassen 1 oder 2 eingestufte "Total Return"- und "Geldmarkt Plus"-Fonds, die im Rahmen der Finanzkrise zu unrühmlicher Bekanntheit gelangt sind. Bei ihnen hat die Finanzkrise offengelegt, dass sie deutlich risikoreicher investiert waren als dies angesichts ihrer niedrigen Risikoeinstufung zu rechtfertigen gewesen wäre. In "guten" Zeiten führte diese Risikoexposition zu einer Überrendite, die als "Alpha" deklariert werden konnte. Im Crash wurden daraus dann allerdings Verluste in einer Größenordnung, die im Bereich eines reinen Aktienfonds liegen.

Die Reaktion des Gesetzgebers auf solche Probleme bestand darin, dass er Anlageberater zum 1. Januar 2010 zur Erstellung eines schriftlichen "Beratungsprotokolls" verpflichtet hat, das die vom BGH bereits 1993 aufgestellten Anforderungen an eine anlage- und objektgerechte Beratung widerspiegelt. Dazu zählt, den Kunden im Rahmen einer Kapitalanlageberatung über alle Eigenschaften und Risiken von Finanzprodukten zu unterrichten, "die für eine getroffene Anlageentscheidung wesentliche Bedeutung haben oder haben können" (BGH v. 6.7.1993, XI ZR 12/93). Ob das Beratungsprotokoll in der Praxis tatsächlich dazu führen wird, dass mehr oder bessere Risikoaufklärung betrieben wird, muss die Zukunft zeigen. Kritiker meinen, dass der einseitige Ruf nach mehr Regulierung die Probleme des Anlegerschutzes nicht lösen kann und verweisen in diesem Zusammenhang beispielsweise darauf, dass auch das Inkrafttreten der ersten Fassung des Wertpapierhandelsgesetzes im Januar 1995 nicht verhindern konnte, dass wenige Jahre später viele Anleger sehr viel Geld beim Platzen der Internetblase verloren haben. Die Kritiker befürchten daher, dass die neue Protokollpflicht letztlich nur dazu führen wird, dass aus Sicht des Privatanlegers die "Papierinflation" im Rahmen eines Beratungsgesprächs zunimmt und in diesem Zuge die Bereitschaft, sich mit den einzelnen Dokumenten tatsächlich auseinanderzusetzen, eher noch weiter sinkt.²

Darüber hinaus ist es wie oben bereits angedeutet so, dass auch die umfangreichste Risikoaufklärung letztlich nutzlos ist, wenn sie auf mangelhaften Methoden der Risikoanalyse und Risikobewertung beruht. Auch diesbezüglich hat die Finanzkrise offengelegt, dass die traditionellen Methoden der Risikomessung erhebliche Schwächen aufweisen. Worin diese im einzelnen bestehen, wird unter anderem in einer wissenschaftlichen Studie im Auftrag des Verbraucherschutzministeriums von Hackethal et al. (2011) ausführlich dargelegt. Die Autoren gelangen dort zu dem Ergebnis, dass eine tatsächlich sachgerechte Risikoaufklärung von Anlegern mit deutlich differenzierteren Modellen und Darstellungen arbeiten müsste, als dies bislang der Fall war, und dass auch nur so "potenzielle Manipulationen" (ebd., S. 115) seitens der Produktanbieter ausgeschlossen werden könnten. Gleichzeitig erscheint den Autoren dieses Ziel allerdings als "nicht realisierbar", da sie bezweifeln, dass "Investoren die so errechneten Risikomaße überhaupt verstehen würden" (ebd.).

Die im Folgenden dargestellten fachlichen Grundlagen, auf denen wir Risikoanalyse und Risikobewertung betreiben, reflektieren unsere eigene, seit Gründung des Instituts im Jahr 2005 stattfindende Auseinandersetzung mit den Schwächen traditioneller Methoden der Risikomessung. Die im Anschluss an diese theoretischen Ausführungen dargestellten Anwendungsbeispiele zeigen, dass wir im Unterschied zu Hackethal et al. (2011) die Realisierbarkeit einer sachgerechten Risikoaufklärung deutlich weniger skeptisch sehen. Unserer Erfahrung nach ist es mit Hilfe geeigneter Darstellungen und Erläuterungen durchaus möglich, Risikoaufklärung in einer für den durchschnittlichen Privatanleger verständlichen Art und Weise zu betreiben, ohne dabei wesentliche fachliche Erfordernisse aufzugeben. Das lässt sich unseres Erachtens beispielsweise anhand unseres schon seit über 7 Jahren betriebenen Verfahrens zur Portfoliozertifizierung zeigen, im Rahmen dessen wir die Verlustpotenziale, die sich in den Jahren 2008 und 2009 realisiert haben, in ihrer Größenordnung durchaus korrekt prognostiziert haben.

²Die Bereitschaft von Privatanlegern, sich in fachlich fundierter Weise mit ihren Kapitalanlagen zu beschäftigen, ist oft tatsächlich erstaunlich gering. Manche Menschen investieren mehr zeitlichen und intellektuellen Aufwand in den Kauf eines Haushaltsgeräts als in die Anlage relevanter Teile ihres Vermögens.

2 Fachliche Grundlagen der Risikoanalyse

2.1 Grundüberlegungen

Jede Entscheidung für oder gegen die Investition in ein Finanzprodukt – sei es ein Sparbuch, ein Mischfonds oder ein Optionsschein – drückt eine Erwartungshaltung über die Zukunft aus und stellt somit eine “Entscheidung unter Unsicherheit bzw. Unwissenheit” (Lazard Asset Management 2009) dar. Dies gilt umso mehr, je weiter diese Erwartungshaltung in die Zukunft reicht, d. h. je länger die geplante Haltedauer für die jeweilige Investition ist. Jede Form des Risikomanagements ist somit lediglich der Versuch, ausgehend von einer mehr oder weniger stark ausdifferenzierten Modellwelt die Zukunft zu prognostizieren. Kein seriöser Risikomanager wird dabei für sich in Anspruch nehmen, sämtliche möglichen Zustände in der Zukunft und deren Zusammenwirken antizipieren zu können. Insofern ist eine der wichtigsten Anforderungen an jedes Prognosemodell, die verwendeten Begriffe, den Geltungsanspruch und somit gleichzeitig die Grenzen des Modells klar zu definieren. Denn wie spätestens die Finanzmarktkrise gezeigt hat, wird es “immer Szenarien geben, mit denen man nicht gerechnet hat” (ebd.), so dass es gilt, “demütig zu bleiben, was die Einschätzung der eigenen Weitsicht angeht” (ebd.).

Bereits der Begriff “Risiko” ist ein bestimmtes Konstrukt innerhalb der Modellwelt des professionellen Risikomanagements und bezeichnet dort nicht die Gesamtheit aller Unsicherheitsfaktoren, denen z. B. ein Wertpapier zukünftig ausgesetzt sein kann, sondern “nur” die benennbaren und quantifizierbaren Formen von Unsicherheit. Die Menge aller möglichen Unsicherheitsfaktoren ist letztlich nie bekannt, denn in extremen Krisensituationen kann es zu Ereignissen kommen, mit denen im Vorfeld praktisch niemand gerechnet hat. Das ist allerdings keine Besonderheit von Investitionen in Wertpapiere, sondern das gilt für jede Form der finanziellen Positionierung. Insofern lässt sich auch nicht sagen, dass Investitionen in Wertpapiere grundsätzlich risikoreicher sind als andere Formen der Investition. Wer zum Beispiel Geld auf ein Sparbuch legt, setzt es dort einem Inflationsrisiko aus, wer es zu Hause lagert, setzt es darüber hinaus einem Diebstahlsrisiko aus. Die Risiken scheinbar risikoloser Formen der finanziellen Positionierung sind oftmals einfach nur weniger offenkundig bzw. schwerer zu quantifizieren.

Durch diese Grundüberlegungen werden bereits einige wichtige Aspekte des Risikobegriffs deutlich:

- Risiko bedeutet im professionellen Risikomanagement ganz allgemein gesprochen die Möglichkeit, dass sich der Wert einer Finanzposition über eine geplante Haltedauer **verändern** kann.
- Dieses Risikoverständnis ist mit dem Alltagsverständnis von Risiko nicht unbedingt identisch: Während im Alltagsverständnis finanzielles Risiko mit der Gefahr von **Verlusten** gleichgesetzt wird, d. h. mit negativen Wertveränderungen, stellt dies im Risikomanagement nur einen Spezialfall dar, der durch sogenannte “Downside“-Maße ausgedrückt wird. Daneben existieren aber auch noch sogenannte “Schwankungs-“ und “Sensitivitätsmaße”, die im Sinne der obigen Definition das gesamte Spektrum der Wertveränderungen im Blick haben, und somit auch positive. Das auch im Kundenkontakt sehr häufig verwendete Risikomaß “Volatilität” ist beispielsweise ein solches Schwankungsmaß, bei dem auch positive Wertveränderungen zum Anstieg führen, was im Sinne des Alltagsverständnis kontraintuitiv ist.
- Finanzielle Risiken unterscheiden sich hinsichtlich ihrer **Erkennbarkeit**: Manche Risiken wie etwa die Kursschwankungen von Aktien sind für jeden offensichtlich, andere wie zum Beispiel das Wiederanlagerisiko von Zinsinstrumenten erschließen sich erst nach näherer Beschäftigung mit der jeweiligen Materie, und wieder andere erschließen sich nur einem kleinen Kreis von Fachleuten, die sich von Berufs wegen mit den jeweiligen Finanzinstrumenten beschäftigen, wie etwa die sog. “Subprimekrise” eindrucksvoll gezeigt hat.³

³Das Beispiel der Subprimekrise zeigt übrigens auch, dass mit einem Risiko natürlich umso eher Missbrauch getrieben werden kann, je schlechter es erkennbar ist und je kleiner somit der Kreis von Personen ist, denen es sich erschließt.

- Finanzielle Risiken unterscheiden sich darüber hinaus hinsichtlich ihrer **Quantifizierbarkeit**: Die Kursrisiken einer permanent gehandelten Aktie sind beispielsweise relativ einfach und präzise zu quantifizieren, während etwa das Ausfallrisiko einer unregelmäßig gehandelten Unternehmensanleihe deutlich schwieriger und nur weniger präzise zu bestimmen ist. Einige Risiken wie zum Beispiel die Stabilität eines politischen oder rechtlichen Umfeldes lassen sich gar nicht oder nur extrem grob quantifizieren. Ein Zusammenhang zwischen der Quantifizierbarkeit und der Bedeutung eines Risikos besteht dabei aber nicht: Ein nicht oder nur schlecht quantifizierbares Risiko kann äußerst wichtig sein und umgekehrt.

2.2 Risikoarten

Wie das Beispiel des Sparbuchs zeigt, müssen verschiedene Arten von Risiken unterschieden werden, die auf unterschiedliche Finanzprodukte unterschiedlich stark einwirken. Hinsichtlich der Risiken von Wertpapieren lassen sich die folgenden Risikoarten unterscheiden:

- Das **Marktrisiko** bezeichnet allgemein das Risiko finanzieller Verluste aufgrund von kurzfristigen Preisänderungen eines Wertpapiers. Das Marktrisiko ist somit in gewisser Weise das "offensichtlichste" Risiko eines Wertpapiers, insbesondere dann, wenn die kurzfristigen Preisschwankungen sehr stark sind. Deshalb wird der Begriff "Marktrisiko" vor allem auf Aktien und andere deutlich volatile Wertpapiere angewendet. Bei einem stark diversifizierten Wertpapierdepot erlaubt die alleinige Betrachtung seines Marktrisikos bereits eine recht gute Beurteilung seines Gesamtrisikos. Bei einzelnen Wertpapieren und insbesondere bei Derivaten bzw. Produkten mit derivativer Komponente ist das aber nicht der Fall, da hier weitere Risikoarten die entscheidende Rolle spielen können.
- Das **Zinsänderungsrisiko** oder Laufzeitenrisiko bezeichnet bei Anleihen und anderen Zinsinstrumenten das Risiko finanzieller Verluste aufgrund von Zinsänderungen. Im Grunde handelt es sich dabei um einen Spezialfall des oben genannten allgemeinen Marktrisikos. Es wird trotzdem als eigene Risikoart betrachtet, da sich Zinsänderungen direkt auf den Marktpreis von Anleihen auswirken, so dass sie für diese wichtige Anlageklasse den zentralen preisbildenden Faktor darstellen.
- Das **Währungsrisiko** bezeichnet das Risiko finanzieller Verluste aufgrund der Abwertung der für einen Investor relevanten Referenzwährung gegenüber der Währung, in der ein vom ihm gehaltenes Wertpapier notiert. Im Grunde handelt es sich auch dabei um einen Spezialfall des allgemeinen Marktrisikos. Es wird deshalb als eigene Risikoart betrachtet, da Wechselkursänderungen für den Halter von Fremdwährungsinstrumenten auch dann ein Risiko darstellen, wenn bei dem jeweiligen Wertpapier in Lokalwährung gar keine Preisänderung stattgefunden hat. Je nach Währungsperspektive kann das Währungsrisiko einer Fremdwährungsposition sogar höher sein als das Marktrisiko des jeweiligen Wertpapiers in Lokalwährung.
- Das **Ausfallrisiko** bezeichnet das für einen Gläubiger bestehende Risiko finanzieller Verluste aufgrund der Unfähigkeit eines Schuldners, seine Verbindlichkeiten vollständig und fristgerecht zu bedienen. Das Ausfallrisiko wird je nach Kontext auch als Emittentenrisiko, Kreditrisiko oder Kontrahentenrisiko bezeichnet. Ausfallrisiken lassen sich einerseits durch qualitative Ratings bewerten, die in einer auf einen Code verdichteten Form Aufschluss über die erwartete Ausfallwahrscheinlichkeit eines Schuldners geben. Andererseits lassen sie sich aber auch durch CDS-Prämien bewerten, die als Versicherungsprämien für eine Versicherung gegen den Ausfall eines Schuldners verstanden werden können. CDS-Prämien besitzen den Vorteil, dass sie direkt am Markt beobachtet werden können und dementsprechend dynamischer auf eine Veränderung der Bonität eines Emittenten reagieren als ein qualitatives Rating. Im Fall von Derivaten dominiert das Ausfallrisiko das Marktrisiko, da die Kursentwicklung des Basiswertes bei einem Ausfall des Emittenten keine Rolle mehr spielt.

- Das **Liquiditätsrisiko** bezeichnet das Risiko finanzieller Verluste aufgrund der Unmöglichkeit, ein Wertpapier kurzfristig am Markt zu liquidieren bzw. nur mit erheblichem Abschlag zu seinem Substanzwert. Gründe hierfür können sein, dass für das angestrebte Geschäft keine Gegenpartei gefunden werden kann, dass der Handel nur unter hohen Geld-Brief-Spannen möglich ist oder dass der Marktpreis des Wertpapiers sehr empfindlich auf getätigte Transaktionen reagiert, so dass beispielsweise bei der Auflösung einer großen Wertpapierposition dessen Marktpreis bereits während der Auflösung fällt. Liquiditätsrisiken sind deutlich schwieriger zu quantifizieren als die oben genannten Risikoarten, weil sie nicht direkt am Markt beobachtet sondern nur indirekt aus bestimmten Marktparametern erschlossen werden können und somit insgesamt weniger genau bestimmbar sind. Ein solcher zumindest indikativer Marktparameter ist etwa die Geld-Brief-Spanne, da diese auf einen Liquiditätsengpass immer mit einer Ausweitung reagieren wird.
- Das **Managementrisiko** bezeichnet bei aktiv gemanagten Finanzprodukten das Risiko finanzieller Verluste aufgrund von ungünstigen Entscheidungen des Managements bei der Positionierung des Wertpapierbestandes relativ zum Markt. Die Risikoprämie für das Tragen des Managementrisikos ist das Alpha. Da die Möglichkeit besteht, positives Alpha zu vereinnahmen, muss für die Partizipation am Managementrisiko im Unterschied zu den oben genannten Risikoarten eine Gebühr gezahlt werden. Diese Gebühr führt dazu, dass das Managementrisiko im Falle von Fehlentscheidungen aus zwei additiv verbundenen Komponenten besteht: Den Verlusten relativ zum Markt plus die gezahlte Gebühr. Insofern ist das Managementrisiko umso höher, je weiter der Handlungsspielraum des Managements ist und je höher die für die Managementleistung geforderte Gebühr ist. Das Managementrisiko lässt sich dadurch auch als eine spezielle Form des Kostenrisikos interpretieren.

Diese Unterscheidung verschiedener Risikoarten führt noch einmal vor Augen, dass es nicht darum gehen kann, alle Arten von Risiken zu vermeiden, sondern darum, sich zu entscheiden, wie stark man welche Risiken eingeht (Jorion 2011, S. 3). Die Vermeidung oder Reduktion eines bestimmten Risikos ist immer nur für den Preis der Erhöhung eines anderen Risikos zu haben. So ist es beispielsweise möglich, dass Marktrisiko eines Indexfonds durch die Investition in einen aktiv gemanagten Fonds zu reduzieren, der durch eine Long-Short-Strategie im selben Markt dessen Marktrisiko reduziert. Diese spezifische Risikoreduktion wird allerdings dadurch erkaufte, dass man sich im Gegenzug einem erhöhten Managementrisiko aussetzt, da die Long-Short-Strategie natürlich funktionieren muss.

2.3 Einfache Risikokennzahlen

Die historisch älteste Kennzahl zur Quantifizierung einer der oben genannten Risikoarten ist die von Frederick R. Macaulay 1938 vorgeschlagene **Duration** zur Messung des Zinsänderungsrisikos. Die nach ihm benannte Macaulay-Duration einer Anleihe ist ein Maß für die Zeit, die ein Anleiheinvestor im Durchschnitt auf Auszahlungen aus dem Wertpapier warten muss. Bezeichnet man die Höhe der Auszahlungen mit c_i und die entsprechenden Zahlungszeitpunkte mit t_i , $1 \leq i \leq n$, so kann der Barwert B der Anleihe berechnet werden als

$$B = \sum_{i=1}^n \frac{c_i}{(1+r)^{t_i}},$$

wobei r die Rendite der Anleihe ist. Die Macaulay-Duration D_{mac} ist dann gegeben durch

$$D_{mac} = \frac{1}{B} \sum_{i=1}^n \frac{t_i c_i}{(1+r)^{t_i}}.$$

Leitet man die oben angegebene Barwertformel nach der Rendite ab, so ergibt sich die sogenannte "Modified Duration", die zur Macaulay-Duration in folgendem unmittelbaren Zusammenhang steht:

$$D_{mod} = -\frac{1}{(1+r)} \cdot D_{mac}.$$

Die Modified Duration wird üblicherweise dazu verwendet, die sog. "Zinselastizität" einer Anleihe zu schätzen, d. h. die Frage zu beantworten, wie stark der Preis einer Anleihe auf eine Zinsänderung reagieren wird. Die Modified Duration kann hier allerdings nur eine erste Schätzung liefern, da sie einen linearen Zusammenhang zwischen dem Preis und der Rendite einer Anleihe unterstellt, obwohl dieser in der Realität konvex ist. Daher wird die Schätzung, die die Modified Duration liefert, umso schlechter sein, je steiler die Zinsstrukturkurve und je größer die Zinsänderung ist. Zu indikativen Zwecken kann die Modified Duration aber oft ausreichend sein, zumal sie aus der Sicht des Halters einer Anleihe unabhängig von der Richtung der Zinsänderung immer zu "pessimistisch" schätzen wird: Bei Zinssenkungen wird sie unterschätzen, wie stark der Preis der Anleihe steigt, und bei Zinsanstiegen wird sie überschätzen, wie stark der Preis der Anleihe fällt. In Zusammenhängen, in denen dieser Schätzfehler nicht toleriert werden kann, muss die Zinselastizität mit Hilfe eines Taylor-Polynoms bestimmt werden, das auch die zweite Ableitung der Barwertformel und somit die Konvexität der Preis-Rendite-Funktion berücksichtigt.

Der Versuch, das Marktrisiko im Sinne des kurzfristigen Preisschwankungsrisikos wissenschaftlich zu quantifizieren, ist historisch fast 15 Jahre jünger und stammt von Harry M. Markowitz (Markowitz 1952) und seinen Kollegen Anfang der 50er Jahre des vergangenen Jahrhunderts. Sie definierten Risiko als die Standardabweichung der Rendite eines Wertpapiers, die kurz auch als die **Volatilität** des Wertpapiers bezeichnet wird. Die Volatilität σ ergibt sich gemäß der folgenden Formel aus der Wurzel der über n Zeitpunkte aufsummierten, quadrierten Abweichungen der Einzelrenditen r_i von ihrem Renditemittelwert r_m :

$$\sigma = \sqrt{\frac{1}{(n-1)} \cdot \sum_{i=1}^n (r_i - r_m)^2}$$

Obwohl die Volatilität bis heute am häufigsten genutzt wird, um das Risiko eines Wertpapiers auszudrücken, ist sie zur Messung des tatsächlichen Anlagerisikos im Sinne der Höhe eines möglichen Verlustes nicht gut geeignet. Das liegt daran, dass es sich bei der Volatilität um ein zweiseitiges Streuungsmaß handelt, das alle Abweichungen vom Mittelwert erfasst, so dass auch außergewöhnlich hohe Gewinne zu einem Anstieg des gemessenen Risikos führen, was kontraintuitiv ist. Bereits Markowitz schlägt deshalb statt der Volatilität die **negative Semivarianz** als Risikomaß vor, die analog zur Volatilität berechnet wird, aber lediglich die negativen Abweichungen von der Durchschnittsrendite berücksichtigt. Dennoch ist auch die Bedeutung der negativen Semivarianz intuitiv nicht leicht verständlich und ermöglicht insbesondere keine Aussagen über die Höhe und Wahrscheinlichkeit von Kapitalverlusten.

Aufbauend auf den Überlegungen von Markowitz haben in den sechziger Jahren William F. Sharpe, John Lintner und Jan Mossin unabhängig voneinander das sogenannte "Capital Asset Pricing Model" (CAPM) entwickelt. Das CAPM ist ein Gleichgewichtsmodell für den Kapitalmarkt und adressiert die Frage, inwieweit sich das gesamte Risiko einer Geldanlage im Sinne der Volatilität durch Diversifikation reduzieren lässt. Grundgedanke des CAPM ist die Zerlegung der erwarteten Rendite eines Wertpapiers μ_i in einen risikolosen Geldmarktzins r_f und eine risikobehaftete Differenz zwischen der Marktrendite μ_m und diesem risikolosen Geldmarktzins. Diese zweite Komponente lässt sich als Risikoprämie interpretieren, die ein Investor dafür erwarten darf, dass er bereit ist, das Marktrisiko bis zu einem bestimmten Grad einzugehen. Entsprechend hängt die zu erwartende Rendite eines Wertpapiers im CAPM davon ab, wie hoch das Exposure des Wertpapiers gegenüber diesen beiden Komponenten ist, so dass die folgende Beziehung gilt (vgl. Sharpe 1964):

$$\mu_i = r_f + \beta_i(\mu_m - r_f)$$

Diese Gleichung macht zunächst deutlich, dass in der Modellwelt des CAPM eine höhere erwartete Rendite immer mit einem höheren Beta, d. h. einem höheren Exposure gegenüber dem Marktrisiko erkaufte werden muss. Eine risikolose Geldanlage mit einem Beta von Null kann entsprechend nur den risikolosen Marktzins erwirtschaften, da andernfalls Marktineffizienzen und Arbitragemöglichkeiten bestehen würden. Die Tatsache wiederum, dass Beta als das Exposure gegenüber dem Marktrisiko interpretiert

werden kann, macht klar, dass das Beta ein Risikomaß darstellt. Im Unterschied zur Volatilität ist es allerdings kein Streuungsmaß, sondern ein Sensitivitätsmaß. Das heißt, es trifft keine Aussage über das absolute Ausmaß von Preisschwankungen, sondern darüber, wie sensibel ein Wertpapier in Relation zu seinem Referenzmarkt auf Preisschwankungen reagiert.

Mathematisch handelt es sich hier um eine Regressionsgleichung, in der das Beta die Steigung der Regressionsgerade ist, für die folgende Beziehung gilt:

$$\beta_i = \frac{\text{Cov}(r_i, r_m)}{\sigma_m^2}$$

wobei r_i die tatsächliche Rendite des Wertpapiers ist, r_m die tatsächliche Rendite des Marktes und σ_m^2 die Varianz des Marktes.

Auch das Beta als Sensitivitätsmaß trifft allerdings keine Aussage über die Höhe und Wahrscheinlichkeit von Kapitalverlusten. Aussagen dieser Art werden durch eine dritte Klasse von Risikomaßen möglich, sogenannte "Downside-Maße", von denen das in der Praxis bekannteste und etablierteste der **Value at Risk** (VaR) ist. Der VaR schätzt den maximal zu erwartenden Verlust, der zum Ende einer vorgegebenen Zeitpanne T mit einer ebenfalls vorgegebenen Sicherheitswahrscheinlichkeit γ nicht überschritten wird. Die Berechnung des VaR ist etwas anspruchsvoller als die Berechnung der Volatilität oder der Semivarianz, weil dafür eine geeignete Verteilungsfunktion unterstellt werden muss, die angibt, mit welchen Wahrscheinlichkeiten welche Wertveränderungen der untersuchten Finanzposition zum Ende der vorgegebenen Zeitspanne zu erwarten sind. Ist eine solche Verteilungsfunktion gegeben, dann entspricht der VaR dem Quantil dieser Wahrscheinlichkeitsverteilung an der Stelle der gewählten Sicherheitswahrscheinlichkeit. Wenn also S_t der Wert einer riskanten Finanzposition ist und $R_t = S_t - S_0$ die Veränderung ihres Wertes bis zum Zeitpunkt t , dann gilt für stetige Verteilungen von S und R :

$$\text{VaR}_\gamma(S) = -F_{R_T}^{-1}(1 - \gamma).$$

Trotz seiner weiten Verbreitung ist allerdings auch der VaR aus mehreren Gründen kein "ideales" Downside-Risikomaß. Das liegt zunächst daran, dass die Lage des Quantils, das zur Bestimmung des VaR herangezogen wird, unabhängig vom Ausmaß der Verluste ist, die unterhalb des Quantils am linken Verteilungsende auftreten können. Dadurch ist es mit Hilfe des VaR nicht möglich, eine Aussage darüber zu treffen, ob bei einer solchen Unterschreitung "nur" mit ungewöhnlich hohen oder mit katastrophal hohen Verlusten gerechnet werden muss. Dieses Defizit des VaR ist beispielsweise bei der Bestimmung von Ausfallrisiken besonders problematisch, da die Ausfallwahrscheinlichkeiten vieler Finanzpositionen unterhalb des typischerweise betrachteten Quantils von 5% liegen und somit trotz ihrer schwerwiegenden Auswirkungen das mit Hilfe des VaR gemessene Risiko nicht beeinflussen.

Eine an dieser Stelle ansetzende Weiterentwicklung des VaR ist der sogenannte **Expected Shortfall** oder **Conditional Value at Risk** (CVaR). Der CVaR ist der bedingte Erwartungswert des Kapitalverlustes, mit dem im Falle einer Unterschreitung des vorgegebenen Quantils zu rechnen ist:

$$\text{CVaR}_\gamma(S) = \mathbb{E}[-R_T | -R_T > \text{VaR}_\gamma(S)].$$

Somit schließt der CVaR das gesamte linke Verteilungsende ein und trifft dadurch eine Aussage zu der oben aufgeworfenen Frage, welche Verluste in jenen besonders ungünstigen Fällen zu erwarten sind, in denen der VaR überschritten ist. Insofern schätzt der CVaR das Verlustrisiko eines Finanzprodukts immer größer ein als der VaR und insofern auch realistischer, als er auch extreme und ungewöhnliche Verteilungsprofile am linken Verteilungsende berücksichtigt. Die folgende Abbildung veranschaulicht dies anhand der Wahrscheinlichkeitsfunktion der Wertveränderungen zweier fiktiver Finanzprodukte I (links) und II (rechts): Während der VaR bei beiden Finanzprodukten identisch ist, weist der CVaR korrekterweise für das Finanzprodukt II ein insgesamt höheres Verlustpotential aus:

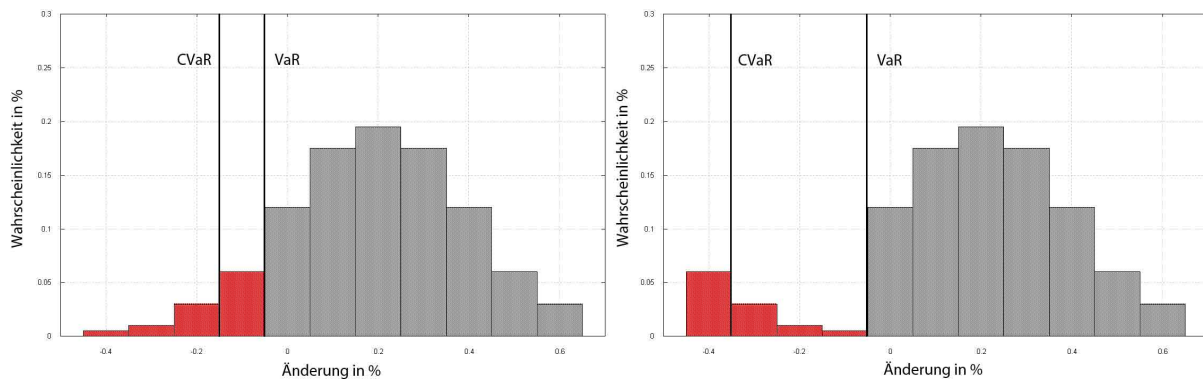


Abbildung 1: VaR und CVaR bei zwei unterschiedlich riskanten Finanzprodukten

Mathematisch gesehen hat der CVaR darüber hinaus den Vorteil, dass er ein sogenanntes **kohärentes Risikomaß** ist, d. h. er erfüllt die Eigenschaften der "Translationsinvarianz", "Subadditivität", "Positiven Homogenität" und "Monotonie", was für die Volatilität nie und für den VaR zumindest bei stark nicht-linearen Auszahlungsprofilen wie im Falle von Optionsscheinen nicht unbedingt gilt. Artzner et al. (1997) illustrieren dies an folgendem Beispiel:

Man stelle sich ein Portfolio vor, welches von zwei Tradern A und B verwaltet wird. Trader A hält einen Put, der weit ausserhalb des Geldes liegt und am nächsten Tag ausläuft, als Shortposition. Trader B hält einen Call, der ebenfalls weit außerhalb des Geldes liegt und am nächsten Tag ausläuft, als Shortposition. Historische VaR-Berechnungen haben ergeben, dass jede der beiden Optionen zum Fälligkeitstermin mit einer Wahrscheinlichkeit von 4% im Geld liegt und dementsprechend einen Verlust produziert. Trader A und B halten also einzeln betrachtet jeweils eine Position, die mit einer Wahrscheinlichkeit von 96% kein Geld verliert, so dass jede der beiden Positionen einen 95%-VaR von Null aufweist. Beim gesamten Portfolio beträgt die Wahrscheinlichkeit für ein verlustfreies Ergebnis jedoch nur 92%, was einen positiven 95%-VaR zur Folge hat. Damit ergibt sich ein Szenario, in welchem das gemeinsame Risiko der beiden Positionen größer ausfällt als die Summe der einzelnen Risiken, d. h. ein negativer Diversifikationsnutzen (Artzner et al. 1997, *eigene Übersetzung*).

2.4 Verteilungsbasierte Risikomodellierung

Es ist deutlich geworden, dass der VaR und insbesondere der CVaR zwei wesentliche Vorteile besitzen: Sie weisen aus mathematischer Sicht bessere Eigenschaften auf als etwa die Volatilität, und sie ermöglichen Aussagen über maximal zu erwartende Verlustpotenziale, was sowohl dem intuitiven Risikoverständnis recht gut entspricht als auch der professionellen Risikosteuerung dient. Aus praktischer Sicht haben diese beiden Risikomaße allerdings den Nachteil, dass sie anspruchsvoller zu bestimmen sind als einfache Risikokennzahlen, da für ihre Berechnung eine Annahme über die Verteilung von Wertpapierrenditen getroffen werden muss, die den Berechnungen in Form einer Wahrscheinlichkeitsverteilung zugrunde gelegt wird. Entsprechend kann die Aussagekraft eines VaR oder CVaR immer nur so gut sein wie die Sachgerechtigkeit der zugrunde gelegten Verteilungsannahme. Somit erhebt sich an dieser Stelle die ganz zentrale Frage, welche Wahrscheinlichkeitsverteilung für die Schätzung der Risikomaße unterstellt wird.

2.4.1 Verteilungsarten

Die einfachste Möglichkeit ist die Verwendung einer **empirischen Verteilung**, die durch historische Daten gegeben ist. Dieses Vorgehen ist allerdings mit zwei wesentlichen Schwierigkeiten verbunden: Zum

einen ist in diesem Fall die Schätzung der Risikomaße ganz wesentlich von der Lage und Länge des verwendeten historischen Zeitintervalls abhängig. Zum anderen ist es dann, wenn der VaR oder CVaR über eine längere Zeitspanne T bestimmt werden sollen (z. B. 1 Jahr), nicht mehr möglich, das benötigte Quantil bzw. bedingte Moment auf Basis eines Zufallsexperiments zu bestimmen, bei dem die Wertveränderung einer Position in voneinander unabhängigen Zeitintervallen beobachtet wird. Wollte man beispielsweise mit einer Sicherheitswahrscheinlichkeit von 95% den CVaR über eine Zeitspanne von einem Jahr anhand historischer Daten bestimmen, dann bräuchte man hierfür mindestens 20 Realisationen unterhalb des unteren 5%-Quantils, da erst davon ausgegangen werden könnte, dass die Daten gegen eine Normalverteilung konvergieren. 20 Realisationen unterhalb des unteren 5%-Quantils würden allerdings 400 unabhängige Realisationen des untersuchten Zufallsexperiments ("Ziehung von 1-Jahres-Renditen") erforderlich machen, d. h. eine Datenhistorie von 400 Jahren. Da eine solche Datenhistorie nicht verfügbar ist, kann in der Praxis eine empirische Verteilung nur dann verwendet werden, wenn der VaR oder CVaR nur für sehr kurze Zeitspannen (1 bis 10 Tage) berechnet werden. Das ist allerdings unbefriedigend bei der Analyse aller Finanzpositionen, die deutlich länger gehalten werden, wie zum Beispiel fast alle Portfolios, die von Privatanlegern für den langfristigen Vermögensaufbau gehalten werden. In solchen Fällen kann man zwar dazu übergehen, rollierende Zeitintervalle der interessierenden Länge zu verwenden, allerdings verletzt man damit die Forderung, dass es sich um unabhängige Realisationen handeln sollte.

Ein grundsätzlich anderer Weg zur Bestimmung von VaR und CVaR ist deshalb die Verwendung einer **theoretischen Verteilungsfunktion** zur Modellierung der Wertpapierrenditen. Die wegen ihrer schönen mathematischen Eigenschaften einfachste Verteilungsfunktion, die man dabei verwenden kann, ist die Normalverteilung bzw. die mit ihr sehr eng verwandte Lognormalverteilung. Dank ihrer Einfachheit war die Lognormalverteilung lange Zeit auch tatsächlich die am häufigsten verwendete Verteilungsfunktion in finanzmathematischen Anwendungen. Spätestens mit Ausbruch der Finanzkrise im Jahr 2008 wurde allerdings klar, dass die (Log-)Normalverteilung zur Abschätzung von Randereignissen, die weit abseits des Mittelwerts liegen, nicht gut geeignet ist, weil sie dazu neigt, die Wahrscheinlichkeit solcher Extremereignisse und somit das tatsächliche Verlustrisiko erheblich zu unterschätzen. Wesentlicher Grund hierfür ist die Tatsache, dass reale Finanzmarktdaten in den meisten Fällen mit einer ausgeprägten "Heteroskedastizität" behaftet sind, was bedeutet, dass die Volatilität über die Zeit nicht konstant ist, sondern mehr oder weniger starken Schwankungen unterworfen ist. Genauer gesagt bilden die täglichen Volatilitäten von Kapitalmarktrenditen sogenannte "Volatility Cluster", was heißt, dass Kurstage mit betragsmäßig hohen Ausschlägen relativ dicht aufeinander folgen. Diese Volatility Cluster haben zur Folge, dass sich bei der empirischen Verteilung von Kapitalmarktrenditen mehr Wahrscheinlichkeitsmasse in den Verteilungsenden konzentriert, als die Normalverteilung abbilden kann (Taleb 2008). Diese wesentliche Schwäche der Normalverteilung illustriert die folgende Abbildung, indem sie die tatsächlichen Wochenrenditen des MSCI Europe (blaue Linie) den per Simulation normalverteilten Renditen mit identischem Erwartungswert und identischer Varianz (rote Linie) gegenüberstellt:



Abbildung 2: Simuliert normalverteilte Renditen im Vergleich zu realen Renditen

Es ist deutlich zu sehen, dass innerhalb der realen Markthistorie insbesondere in 2008 Renditen auftreten, die bei Unterstellung einer Normalverteilung nicht erklärbar sind. Eine Wochenrendite unterhalb von -20%, wie sie beim MSCI Europe zwischen dem 3. und dem 10. Oktober 2008 tatsächlich auftrat, liegt unter Normalverteilungsannahmen über sieben Standardabweichungen vom Mittelwert entfernt. Eine so extreme Wochenrendite würde die Normalverteilung lediglich mit einer Wahrscheinlichkeit von rund $6,6 \cdot 10^{-13}$, oder anders ausgedrückt einmal in 30 Milliarden Jahren erwarten.

Um diesem Problem zu begegnen, lassen sich nun zwei verschiedene Wege einschlagen: Die erste Möglichkeit besteht darin, eine Verteilungsfunktion zu verwenden, die mehr Wahrscheinlichkeitsmasse in den Verteilungsenden besitzt und somit extreme Ereignisse mit einer höheren Wahrscheinlichkeit erwartet. Diese Eigenschaft einer Verteilung wird allgemein als **Fat Tails** bezeichnet. Fat Tails lassen sich beispielsweise durch die regimeabhängige Mischung von zwei oder mehreren Verteilungen modellieren, wobei zur Mischung durchaus auch Normalverteilungen verwendet werden können.⁴ Weitere geeignete Verteilungen können etwa die t -Verteilung oder die Exponential-Power-Verteilung⁵ sein, die beide auch unter praktischen Gesichtspunkten noch gut handhabbar sind. Allerdings erlaubt die t -Verteilung die Modellierung von Fat Tails lediglich bis zu einer Exzess-Wölbung von sechs bei fünf Freiheitsgraden – für kleinere Freiheitsgrade ist die Wölbung nicht definiert. Die Exponential-Power-Verteilung lässt sich aus der Normalverteilung ableiten, wenn man bei der Dichte der Normalverteilung den Ausdruck in der Exponentialfunktion nicht mehr mit 2, sondern mit einem variablen Parameter $\beta > 0$ potenziert. Im Fall $\beta = 2$ resultiert die klassische Normalverteilung, und im Fall $\beta < 2$ besitzt die Verteilung Fat Tails. Die Exponential-Power-Verteilung besitzt allerdings den Nachteil, dass sie für kleine Exponenten ihre Glockenform verliert und beim Auftreten von Fat Tails eine ausgeprägte Spitze besitzt.

Die zweite Möglichkeit, der Fat Tails Problematik zu begegnen, besteht darin, eine sogenannte **Monte-Carlo Simulation** durchzuführen. In diesem Fall wird die Verteilung der Wertentwicklung eines Finanzinstruments über den interessierenden Zeitraum nicht mehr mit Hilfe einer a priori zugrunde gelegten theoretischen Verteilung analytisch berechnet, sondern es wird ein stochastischer Prozess unterstellt, der Wertveränderungen bewirkt, so dass sich durch dessen häufige Wiederholung Kurspfade simulieren lassen. Durch eine große Anzahl von simulierten Kurspfaden kann dann wiederum eine Verteilung approximiert werden. Auf diese Weise können deutlich flexiblere Wahrscheinlichkeitsverteilungen generiert werden und auch eventuelle zeitliche Abhängigkeiten zwischen den simulierten Renditen ohne die relativ komplizierte n -fache Faltung einer Wahrscheinlichkeitsverteilung abgebildet werden. Dies gilt aber natürlich nur dann, wenn der unterstellte stochastische Prozess ein sachgerechtes Modell der Wertveränderungsdynamik für das jeweilige Finanzinstrument darstellt. Bei der Betrachtung von Aktienmarktrenditen hat sich hier insbesondere das GARCH-Modell bewährt, da es die in den Daten üblicherweise vorliegende Heteroskedastizität, Volatility Clusters und Fat Tails in aller Regel ausreichend berücksichtigt. Es wird daher im Folgenden näher vorgestellt.

2.4.2 Simulation von Aktienmarktrenditen

Der Grundgedanke des **GARCH-Modells**⁶ besteht darin, angesichts der Existenz von Volatility Clusters die Volatilität nicht mehr als konstanten Parameterwert anzusehen, sondern sie selbst als Zeitreihe zu modellieren, die durch vergangene Kursbewegungen beeinflusst wird. Dabei wird die aktuelle Varianz auf Basis der zuletzt vorliegenden Varianzen und den zuletzt aufgetretenen quadrierten Fehlertermen bestimmt (vgl. Bollerslev 1986). Damit entspricht das GARCH-Modell dem für die Modellierung des Zeitreihenmittelwerts gebräuchlichen ARMA-Modell und besitzt die folgende Gestalt:

$$\sigma_t^2 = \omega + \alpha \sigma_{t-1}^2 + \beta \varepsilon_{t-1}^2$$

Die obenstehende Formel stellt den GARCH-Prozess erster Ordnung dar, d. h. es werden jeweils nur die zuletzt gültige Varianz und der zuletzt beobachtete Fehlerterm in die Gleichung für die aktuelle Varianz

⁴Insofern ist es auch falsch zu behaupten, dass mit Hilfe von Normalverteilungen grundsätzlich keine Fat Tails modelliert werden können.

⁵Die Exponential-Power-Verteilung wird auch als "Verallgemeinerte Normalverteilung" bezeichnet.

⁶Generalized Autoregressive Conditional Heteroskedasticity

aufgenommen. Für die Praxis ist diese einfachste und im Hinblick auf die Anzahl der zu schätzenden Parameter sparsamste Version des GARCH-Modells fast immer ausreichend.

Da die wöchentlichen Kapitalmarktrenditen innerhalb des GARCH-Modells nicht mehr unabhängig sind, sondern rekursiv voneinander abhängen, kann die Wahrscheinlichkeitsverteilung für die Wertentwicklung des Marktes über einen längeren Zeitraum (z. B. 1 Jahr) nur noch durch die Monte-Carlo Simulation einer großen Anzahl von Kurspfaden praktikabel geschätzt werden. Das bedeutet, dass die n -te Wochenrendite erst dann simuliert werden kann, wenn die $(n - 1)$ -te Wochenrendite bereits beobachtet bzw. simuliert wurde, da die für den n -ten Simulationsschritt benötigten Verteilungsparameter von der $(n - 1)$ -ten Realisation abhängen.

Neben der Fat Tails Eigenschaft kann die Wahrscheinlichkeitsverteilung von Finanzmarktrenditen weiterhin eine gewisse Asymmetrie oder Schiefe aufweisen, wobei diese Eigenschaft innerhalb der Aktien- und Rohstoffmärkte in der Regel nicht so stark ausgeprägt ist wie die allgegenwärtigen Fat Tails. Im Folgenden wird deshalb eine einfache Möglichkeit dargestellt, wie sich durch die Mischung von zwei Normalverteilungen auf weitgehend unkomplizierte Weise eine schiefe Verteilung konstruieren lässt.

Die Mischung von zwei Normalverteilungsdichten $f_1(x) = f(x, \mu_1, \sigma_1^2)$ und $f_2(x) = f(x, \mu_2, \sigma_2^2)$ anhand eines Parameters $0 \leq \lambda \leq 1$ resultiert in einer relativ flexiblen Mischverteilung mit fünf Parametern (vgl. Abbildung 3):

$$f(x, \mu_1, \sigma_1^2, \mu_2, \sigma_2^2, \lambda) = \lambda f(x, \mu_1, \sigma_1^2) + (1 - \lambda) f(x, \mu_2, \sigma_2^2)$$

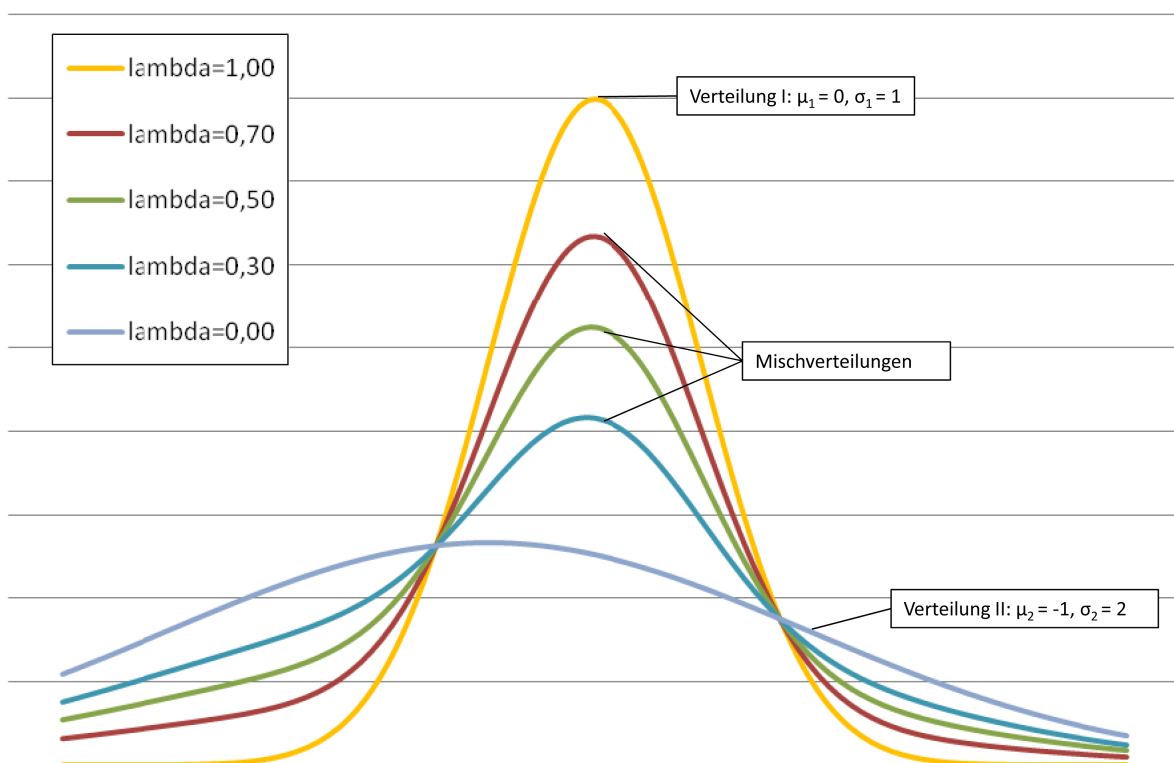


Abbildung 3: Mischverteilungen aus zwei Normalverteilungen

Der wesentliche Vorteil dieser Klasse von Verteilungen liegt in der im Vergleich zu anderen schiefen Verteilungen relativ einfachen Handhabung, da das vorgeschlagene Modell auf der Normalverteilung aufbaut und sich dementsprechend relativ leicht anpassen und simulieren lässt.

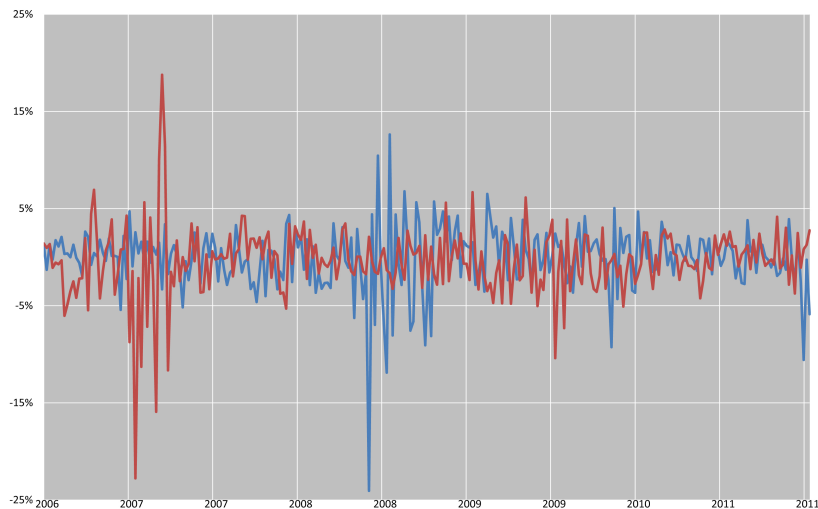


Abbildung 4: GARCH-simulierte Renditen im Vergleich zu realen Renditen

Wie Abbildung 4 zu entnehmen ist, ermöglicht das GARCH-Modell in Verbindung mit gemischt normalverteilten Residuen insgesamt eine angemessene Abbildung der realen Wochenrenditen. Auch innerhalb eines Kolmogorow-Smirnow-Tests kann die Nullhypothese, d. h. die Verteilungsgleichheit der beiden dargestellten Zeitreihen, mit einem p-Wert von 12,1% zu keinem in der Praxis gebräuchlichen Signifikanzniveau abgelehnt werden.

Innerhalb des GARCH-Modells müssen die drei Parameter ω , α und β geschätzt werden. Die unbedingte Varianz innerhalb des GARCH-Modells beträgt

$$\sigma^2 = \frac{\omega}{1 - \alpha - \beta},$$

so dass der Parameter ω durch

$$\hat{\omega} = (1 - \hat{\alpha} - \hat{\beta}) S^2$$

geschätzt werden kann, um eine unverzerrte Modellierung der Varianz zu ermöglichen. Die verbleibenden Parameter α und β werden mit Hilfe eines **Quasi-Maximum-Likelihood (QML)** Verfahrens geschätzt, was bedeutet, dass für die Berechnung und Optimierung der Likelihood-Funktion eine fehlspezifizierte Verteilung zu Grunde gelegt wird, auf Grund ihrer einfachen Handhabung häufig eine Normalverteilung. QML-Schätzer sind unter gewissen Regularitätsbedingungen konsistent und asymptotisch normal. Für die Parameter α und β müssen die folgenden Nebenbedingungen erfüllt sein:

1. $\alpha, \beta \geq 0$
2. $\alpha + \beta \leq 1$ muss darüber hinaus für einen stationären Varianzprozess gelten.

Die (Quasi-)Likelihood-Funktion besitzt dann die folgende Gestalt:

$$L(\alpha, \beta) = L(\alpha, \beta | \varepsilon_1, \dots, \varepsilon_T) = \prod_{t=1}^T f_{\alpha, \beta}(\varepsilon_t)$$

wobei $\varepsilon_t = r_t - \mu$ das t-te Residuum und $f_{\alpha, \beta}(\varepsilon_t)$ die Dichte der Normalverteilung mit Erwartungswert 0 und Varianz ($t > 1$)

$$\sigma_t^2 = (1 - \alpha - \beta) S^2 + \alpha \sigma_{t-1}^2 + \beta \varepsilon_{t-1}^2$$

ist. Im Fall der ersten Beobachtung ($t = 1$) kann die unbedingte Varianz bzw. deren Schätzung in Gestalt der Stichprobenvarianz S^2 als Startwert verwendet werden. Da das Produkt über die einzelnen

Dichtefunktionen schwierig zu maximieren ist, wird die Likelihood-Funktion in praktisch allen Fällen noch logarithmiert, was wegen der Monotonie des Logarithmus keinen Einfluss auf das optimale Ergebnis hat.

$$\log(L(\alpha, \beta)) = \sum_{t=1}^T \log(f_{\alpha, \beta}(\varepsilon_t))$$

Der Maximum-Likelihood Schätzer ergibt sich dann durch diejenigen Parameterwerte $\hat{\alpha}$ und $\hat{\beta}$, welche die Log-Likelihood Funktion maximieren.

$$(\hat{\alpha}, \hat{\beta}) = \arg \max_{\alpha, \beta} \log(L(\alpha, \beta))$$

Die Schätzung der fünf Parameter $\mu_1, \sigma_1, \mu_2, \sigma_2$ und λ der gemischten Normalverteilung, die als bedingte Verteilung der standardisierten Residuen verwendet wird, funktioniert weitgehend analog wie im Falle des GARCH-Modells, mit dem Unterschied, dass diesmal auf Basis der echten Likelihood-Funktion geschätzt wird. Dazu müssen die Residuen zunächst an Hand der korrespondierenden Varianzen, die zuvor mit Hilfe des GARCH-Prozesses modelliert wurden, standardisiert werden.

$$\epsilon_t = \frac{\varepsilon_t}{\sigma_t}$$

Anschließend wird die Log-Likelihood auf Basis der Mischverteilung gebildet.

$$\log(L(\mu_1, \sigma_1, \mu_2, \sigma_2, \lambda)) = \sum_{t=1}^T \log(\lambda f_{\mu_1, \sigma_1}(\epsilon_t) + (1 - \lambda) f_{\mu_2, \sigma_2}(\epsilon_t))$$

Durch Maximierung der Log-Likelihood lassen sich wiederum die optimalen Parameterwerte $\hat{\mu}_1, \hat{\sigma}_1, \hat{\mu}_2, \hat{\sigma}_2$ und $\hat{\lambda}$ feststellen.

2.4.3 Simulation von Zinszeitreihen

Die Wertveränderungsdynamik von Anleihekursen weist wesentliche Unterschiede zu derjenigen von Aktien- oder auch Wechselkursen auf: So zeigt der Kurs einer Anleihe eine sogenannte "Mean Reversion", d. h. er strebt über die Zeit gegen den Nennwert der Anleihe, zu dem sie nach Ablauf zurückgezahlt wird (sofern die Anleihe zuvor nicht ausgefallen ist). Damit zusammenhängend weist die Volatilität von Anleihekursen ebenfalls eine zeitabhängige Tendenz auf, da sie über die Zeit gegen den am Rückzahlungstag erreichten Wert "Null" strebt. Desweiteren können (nominale) Zinssätze in aller Regel keine negativen Werte annehmen. Aus diesen Gründen muss für die Simulation von Zinszeitreihen ein stochastischer Prozess zugrunde gelegt werden, der diese Sachverhalte berücksichtigt. Hierfür sind eine ganze Reihe von Zinsmodellen vorgeschlagen worden, die sich im Wesentlichen in zwei Gruppen einteilen lassen: Die sogenannten "Equilibrium-" oder "Short-Rate-Modelle" unterstellen einen stochastischen Prozess für die Entwicklung des Zinssatzes über einen theoretisch infinitesimal kurzen Zeitraum (die Short Rate), und modellieren darüber, ausgehend von der aktuellen Short Rate, die Zinskurve (Cox et al. 1985). Demgegenüber unterstellen sogenannte "No Arbitrage"- oder "Forward-Rate-Modelle" einen stochastischen Prozess für die Entwicklung der gesamten Zinskurve (Ho and Lee 1986, Hull and White 1990, Heath et al. 1992).

Ein eher einfaches, aber für unsere Zwecke in der Regel ausreichendes Zinsmodell aus der Gruppe der Short-Rate-Modelle ist das Modell von **Cox, Ingersoll und Ross (CIR)**, welches einerseits Mean-Reversion enthält und den zufälligen Term andererseits mit der Wurzel aus dem aktuellen Zinssatz gewichtet, so dass der zufällige Term bei sehr niedrigen Zinssätzen weitgehend verschwindet und der deterministische Term eine Rückkehr zum Mittelwert forciert (vgl. Cox et al. 1985).

$$dr_t = \theta(\mu - r_t) dt + \sigma\sqrt{r_t} dW_t$$

W_t bezeichnet dabei den Wiener-Prozess, d. h. einen stetigen Prozess mit unabhängigen, stationären und normalverteilten Zuwächsen. Damit das Simulationsmodell in den (sehr seltenen) Fällen eines negativen Zinssatzes nicht zu einem Fehler führt, wird die Wurzelfunktion in den negativen Bereich erweitert.

$$f(x) = \operatorname{sgn}(x) \sqrt{|x|}$$

Insgesamt lassen sich mit dem CIR-Prozess über den betrachteten Zeithorizont von einem Jahr hinweg realistische Zinszeitreihen simulieren (vgl. Abbildungen 5 und 6).

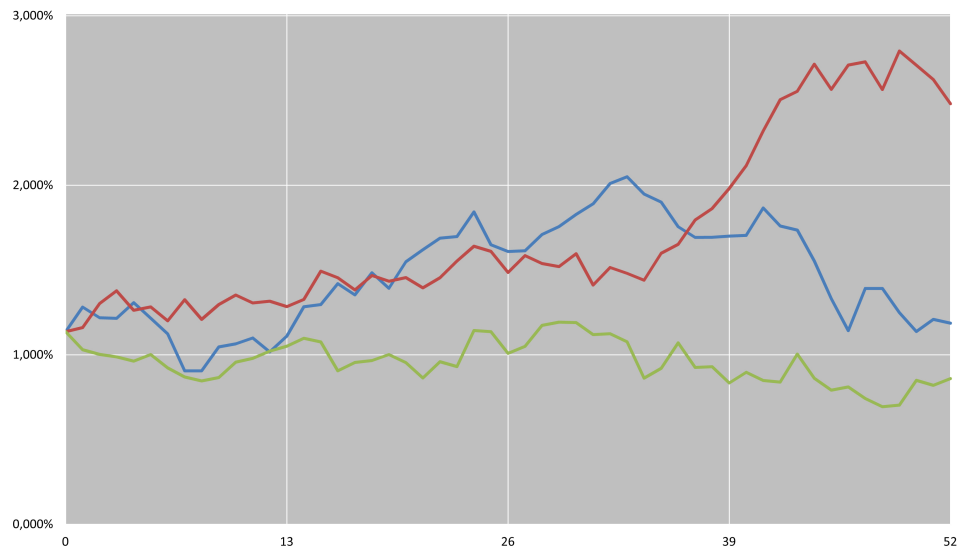


Abbildung 5: Simulierte Pfade einer Zero Rate aus der Zinsstrukturkurve für Bundesanleihen

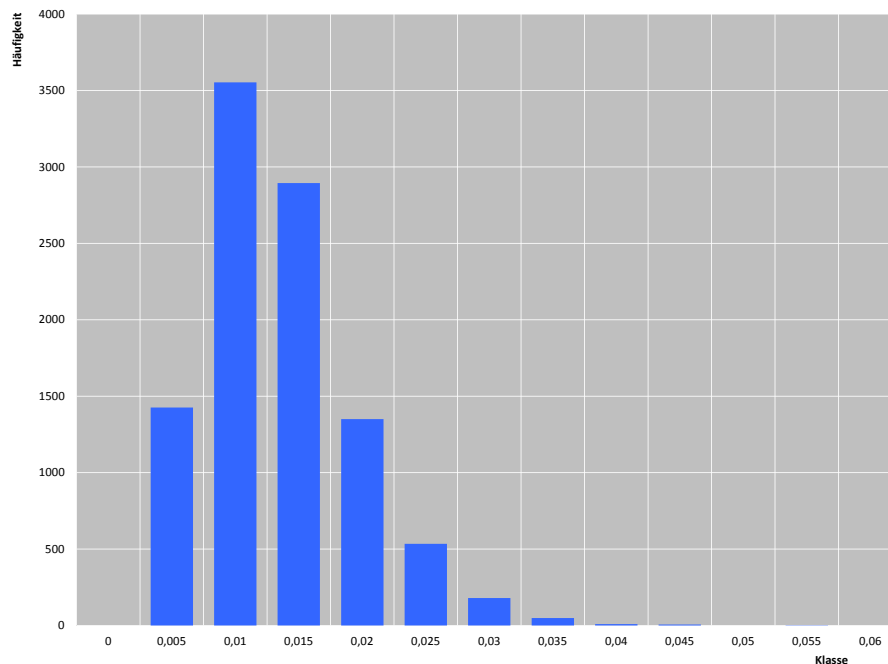


Abbildung 6: Simulierte Verteilung einer Zero Rate aus der Zinsstrukturkurve für Bundesanleihen nach einem Jahr

Die Parameter des CIR-Modells θ , μ und σ werden in mehreren Schritten auf Basis der **Kleinsten Quadrate (KQ)** geschätzt. Im ersten Schritt wird das mittlere Zinsniveau μ über die wöchentlichen Zinssätze der vergangenen fünf Jahre geschätzt.

$$\hat{\mu} = \frac{1}{T} \sum_{t=1}^T r_t$$

Im nächsten Schritt müssen die wöchentlichen Veränderungen des betrachteten Zinssatzes gebildet werden.

$$\Delta r_t = r_{t+1} - r_t$$

Für diese Differenzen gibt es innerhalb des CIR-Modells einen korrespondierenden Schätzwert in Gestalt von:

$$\Delta \hat{r}_t = \theta (\hat{\mu} - r_t)$$

Der Parameter θ kann dementsprechend durch Minimierung der Fehlerquadratsumme

$$Q(\theta) = \sum_{t=1}^{T-1} (\Delta r_t - \theta (\hat{\mu} - r_t))^2$$

bestimmt werden. Der Parameter σ wird im letzten Schritt durch die Stichprobenstandardabweichung der Residuen

$$\hat{\varepsilon}_t = \frac{(\Delta r_t - \hat{\theta}(\hat{\mu} - r_t))}{\sqrt{r_t}}$$

bestimmt.

2.4.4 Berücksichtigung von Ausfallwahrscheinlichkeiten

Wie deutlich geworden ist, besteht ein wesentlicher Unterschied zwischen verschiedenen Verteilungsmodellen darin, wie realistisch sie sehr große Wertveränderungen einer Finanzposition und somit auch sehr hohe Verluste abbilden können. Selbst die in dieser Hinsicht "guten" Verteilungsmodelle treffen allerdings nur so lange gute Aussagen, wie es nicht zu einem Ausfall der Finanzposition und somit zu einem Totalverlust kommt. Bei Finanzinstrumenten, bei denen diese Möglichkeit nicht vernachlässigt werden darf, muss diese Ausfallwahrscheinlichkeit somit explizit berücksichtigt werden. Dies kann mit Hilfe der Prämien von **Credit Default Swaps (CDS)** geschehen.

Ein CDS ist eine Versicherung gegen das Ausfallereignis eines Wertpapieremittenten in einem bestimmten Zeitraum, wobei der Käufer des CDS im Falle eines Ausfalls eine Entschädigungszahlung entsprechend der Recovery Rate RR und des Nominals N erhält. In aller Regel werden CDS auf Senior Tranchen mit einer unterstellten Recovery Rate von 40%⁷ emittiert, d. h. die Höhe der Entschädigungszahlung resultierend aus dem CDS beträgt 60%. Für diese Versicherung müssen Prämienzahlungen entrichtet werden, die als CDS-Spread s in Basispunkten bezogen auf das versicherte Nominal quotiert werden. Die einzelnen Zahlungen finden in der Regel vierteljährlich statt.

Sofern zu bestimmten festgelegten Zeitpunkten $0 < t_1 < \dots < t_n = T$ Prämienzahlungen stattfinden, und der (aktuell unbekannte) Zeitpunkt des Ausfalls des Emittenten durch eine Zufallsvariable τ modelliert wird, so lässt sich unter der Annahme eines risikolosen Zinssatzes r_f der als EDPL (expected discounted premium leg) bezeichnete Barwert der zukünftigen Prämienzahlungen angeben. Unter der Annahme, dass der zufällige Ausfallzeitpunkt τ gemäß einer Exponentialverteilung⁸ mit Parameter λ verteilt ist, gilt:

⁷Im Falle von "Subordinated" Tranchen 20%.

⁸Die Exponentialverteilung ist das klassische Modell für Wartezeiten und Lebensdauern. Im Kontext von Ausfallwahrscheinlichkeiten bedeutet dies, dass durch die Zufallsvariable τ die Wartezeit bis zum Ausfallereignis modelliert wird.

$$\text{EDPL} = \mathbb{E} \left[\sum_{k=1}^n s \Delta t_k e^{-r_f t_k} \mathbb{1}_{\{\tau > t_k\}} \right] = \sum_{k=1}^n s \Delta t_k e^{-(r_f + \lambda) t_k}.$$

Die Schwierigkeit bei der Berechnung des Barwertes der Entschädigungszahlung besteht darin, dass der Emittent zu jedem Zeitpunkt ausfallen kann. Durch die oben getroffene Verteilungsannahme kann das EDDL (expected discounted default leg) berechnet werden durch

$$\text{EDDL} = \mathbb{E} \left[(1 - RR) e^{-r_f \tau} \mathbb{1}_{\{0 \leq \tau \leq T\}} \right] = (1 - RR) \frac{\lambda}{r_f + \lambda} \left(1 - e^{-(r_f + \lambda) T} \right).$$

Der Preis eines CDS wird durch den (eindeutigen) CDS-Spread ausgedrückt, welcher so gewählt wird, dass sich die Werte beider Legs zum heutigen Zeitpunkt ausgleichen. CDS-Spreads sind für viele Unternehmen beobachtbar, die Fremdkapital am Kapitalmarkt aufnehmen, so dass folgender Ansatz zur Ermittlung einer Ausfallwahrscheinlichkeit angewendet werden kann: Kennt man den aktuell vom Markt verlangten CDS-Spread, so kann man die Ausfall- bzw. Überlebenswahrscheinlichkeit ermitteln, mit welcher der Markt für den jeweiligen Emittenten rechnet. Im obigen Modell bedeutet dies, dass λ unter bekanntem s so ermittelt werden muss, dass

$$\text{EDDL} = \text{EDPL}.$$

Hat man λ ermittelt, so kann die Ausfallwahrscheinlichkeit p innerhalb des Zeitraums $[0, t]$ ermittelt werden als

$$p = 1 - e^{-\lambda t}.$$

Dieses Vorgehen ist das Standardvorgehen der International Swaps and Derivatives Association (ISDA) zur Ermittlung von Ausfallwahrscheinlichkeiten, das auf alle Marktteilnehmer mit liquide gehandelten CDS angewendet werden kann, was in der Regel für Staaten und Blue-Chip-Unternehmen der Fall ist.

Um ein allgemein anwendbares Modell zur Ermittlung von Ausfallwahrscheinlichkeiten zu erhalten, lässt sich ein Regressionsmodell verwenden, das den Credit Spread und die Macaulay-Duration D_{mac} einer Anleihe in direkte Verbindung mit der jeweiligen Ausfallwahrscheinlichkeit bringt.

Der Credit Spread gibt die Überrendite zum entsprechenden risikolosen Zins an und ist somit ein Maß für das vom Anleihekäufer akzeptierte Risiko. Wird der Credit Spread als maßgeblich beeinflussender Faktor mit x bezeichnet, so kann das Modell für Ausfallwahrscheinlichkeiten unter Zuhilfenahme der Macaulay-Duration D_{mac} angegeben werden als

$$\hat{p} = a_1 x^2 + a_2 x + a_3 D_{mac} + a_4 x D_{mac} + b,$$

wobei $a_1 \in \mathbb{R}$, $a_2 \in \mathbb{R}$, $a_3 \in \mathbb{R}$, $a_4 \in \mathbb{R}$ und $b \in \mathbb{R}$ die Modellparameter sind und $\hat{p}$ als abhängige Variable die entsprechende Ausfallwahrscheinlichkeit darstellt. Um das Modell anwenden zu können, benötigt man also lediglich fünf Parameter, die auf Basis einer Stichprobe geschätzt werden können, die Informationen zu Ausfallwahrscheinlichkeiten von Emittenten enthält. Konkret schätzt man das Modell, indem man zunächst mit der oben vorgestellten Methodik eine Stichprobe für Emittenten mit gehandelten CDS konstruiert, die als Basis für die Parameterschätzung dient. Bei der Schätzung ist stets darauf zu achten, dass sich das Bestimmtheitsmaß für die Regression in einem akzeptablen Intervall befindet (vgl. Abbildung 7).

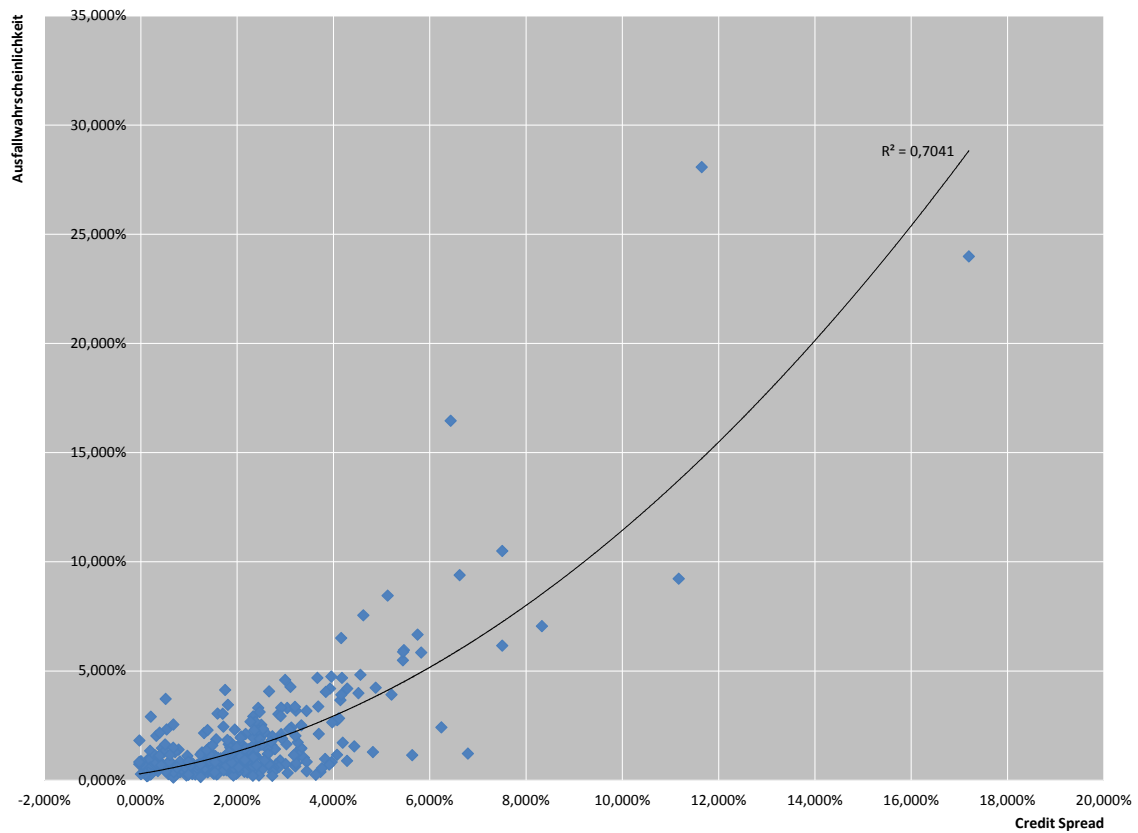


Abbildung 7: Regressionsmodell zur Schätzung von Ausfallwahrscheinlichkeiten bei fehlender CDS-Prämie

3 Bewertungsprozesse

3.1 Klassifizierung von Wertpapieren anhand ihrer Risikomerkmale

Will man nun das Risiko eines konkreten Wertpapiers bewerten, dann ist es notwendig, zunächst dessen prinzipielle Wertveränderungsdynamik zu verstehen. Denn wie die bisherigen Ausführungen gezeigt haben, ist es beispielsweise wichtig zu wissen, ob das Wertpapier relativ leicht ausfallen kann oder ob seine Kurse eine zeitabhängige "Mean Reversion" aufweisen. Nimmt man eine entsprechende Klassifizierung von Wertpapieren anhand ihrer Risikomerkmale vor, so erscheinen uns folgende Unterscheidungen sinnvoll (vgl. Abbildung 8):

Sammelvermögen vs. Einzeltitel:

Dadurch, dass Investmentfonds bzw. allgemein Sammelvermögen das investierte Vermögen auf eine Vielzahl von Einzeltiteln verteilen und dabei in der Regel auch über mehrere Risikofaktoren streuen, sind diese Finanzprodukte vor allem systematischen Risiken ausgesetzt, d. h. Markt-, Zinsänderungs- und Währungsrisiken. Das Ausfallrisiko eines Einzeltitels spielt somit nur eine untergeordnete Rolle und kann bereits in der Datenhistorie des betrachteten Investmentfonds enthalten sein. Es ist deshalb in der Regel hinreichend, bei der Berechnung von Risikomaßen direkt auf der Zeitreihe des betrachteten Investmentfonds aufzubauen.

Demgegenüber ist das Ausfallrisiko bei der Betrachtung eines einzelnen Wertpapiers von erheblicher Bedeutung für das gemessene Gesamtrisiko und muss deshalb zusätzlich in die Risikobewertung aufgenommen werden, da innerhalb der Datenhistorie eines noch existierenden Wertpapiers in der Regel kein Ausfallereignis stattgefunden haben kann.

Festverzinsliche Wertpapiere vs. Wertpapiere mit Aktienrisiko:

Aufgrund der oben genannten Spezifika von Anleihekursen (Mean Reversion, zeitabhängige Volatilitätstendenz, keine negativen Zinsen) muss bei der Bewertung festverzinslicher Wertpapiere ein anderes Simulationsmodell zugrunde gelegt werden als im Falle von Aktien.

Direktinvestitionen vs. Derivate:

Derivate unterscheiden sich von Direktinvestments vor allem darin, dass der Zusammenhang, der zwischen ihrem Rückzahlungsbetrag und der Wertentwicklung ihres Basiswertes besteht, in der Regel nicht linear ist. Vor diesem Hintergrund ist es kaum möglich, die Wertentwicklung eines Derivats direkt zu simulieren. Es ist jedoch in der Regel möglich, den Basiswert zu simulieren und die korrespondierende Wertentwicklung des Derivats aus der simulierten Wertentwicklung des Basiswertes abzuleiten.

Darüber hinaus sind Derivate immer mit einem erhöhten Ausfallrisiko behaftet, in manchen Fällen sogar mit mehreren, die sogar weitgehend unabhängig voneinander sein können – so kann etwa bei einem Zertifikat auf eine Aktie sowohl der Emittent des Zertifikats als auch der zu Grunde liegende Basiswert in Gestalt der Aktie eines Unternehmens insolvent werden, wobei in beiden Fällen mit einem Totalverlust zu rechnen ist.

Simulationsbasierte vs. direkte Zuordnung:

Einige Wertpapiergattungen können aufgrund ihrer spezifischen Risikomerkmale ohne die Berechnung eines quantitativen Risikomaßes direkt einer Risikoklasse zugeordnet werden. Dies gilt beispielsweise für alle Hebelprodukte, bei denen unabhängig von einem Totalverlust des Basiswertes ein signifikantes Totalverlustrisiko besteht. Diese Produkte können direkt der höchsten Risikoklasse zugeordnet werden. Auf der anderen Seite können festverzinsliche Produkte ohne Währungs- und Ausfallrisiko, darunter beispielsweise euronotierte Konten und Sparbriefe, die dem Einlagensicherungsfonds unterliegen, direkt der niedrigsten Risikoklasse zugeordnet werden.

Für die meisten Wertpapiere wird eine solche einfache, direkte Zuordnung allerdings nicht möglich sein, so dass für sie eine quantitative Risikobewertung stattfinden muss. Im Folgenden wird dargestellt, wie wir diese Bewertung im Falle von Investmentfonds und einzelnen Anleihen, Aktien und Derivaten vornehmen.

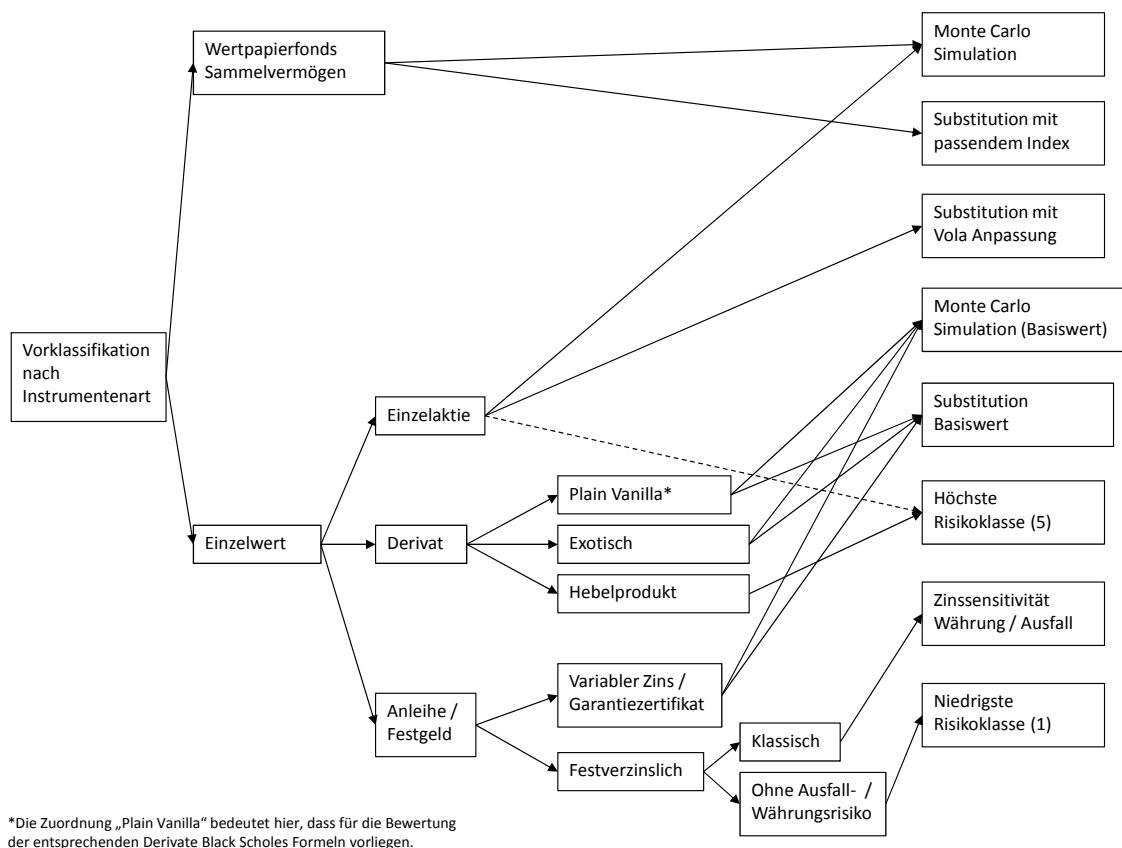


Abbildung 8: Übersicht zur Klassifizierung der einzelnen Wertpapierarten

3.2 Bewertung von Wertpapierfonds / Sammelvermögen

Im Wesentlichen wird das Risikomaß dieser Produktgattung vom Marktrisiko beeinflusst. Die Modellierung des Verlustrisikos erfolgt via Monte-Carlo Simulation, wobei die Renditen des Wertpapiers gemäß einem GARCH-Modell mit gemischter Normalverteilung verteilt sind. Die Modellanpassung erfolgt auf Basis einer Zeitreihe des Wertpapiers in EUR, welche mindestens wöchentliche Kurswerte über die letzten fünf Jahre enthält. Falls die Zeitreihe des Wertpapiers die geforderte Länge nicht aufweist, muss ein geeignetes Substitut zur Schätzung der Modellparameter verwendet werden. Die zu beachtenden Risiken in diesem Modell beschränken sich damit auf das Markt- und das Währungsrisiko. Im einzelnen bedeutet das:

- Das Marktrisiko wird über ein GARCH-Modell simuliert.
- Ein eventuell vorliegendes Währungsrisiko wird bereits in der Zeitreihe des jeweiligen Wertpapiers berücksichtigt, weil die Modellanpassung auf Basis von Marktdaten in EUR erfolgt, weshalb eine gesonderte Betrachtung in diesem Fall nicht mehr erforderlich ist.
- Ausfallrisiken finden lediglich dahingehend Berücksichtigung, dass Ausfallereignisse, die in der Historie des betrachteten Produktes enthalten sind, in die geschätzten Modellparameter eingehen.
- Auf eine quantitative Verarbeitung des Liquiditätsrisikos wird verzichtet. Es wird davon ausgegangen, dass Fondsanteile unter normalen Marktbedingungen an die KAG zurückgegeben werden können. Zwar wäre es denkbar, dass die Rücknahme von Fondsanteilen in einer zukünftigen Finanzkrise ausgesetzt wird, allerdings ist dies mit Ausnahme von Immobilienfonds, die innerhalb des Bewertungsprozesses ohnehin einen Sonderfall darstellen (s. u.), bisher noch nicht in relevantem Umfang vorgekommen.
- Die Preishistorie von Immobilienfonds ist durch die vergleichsweise seltenen Bewertungstermine des Immobilienportfolios stark geglättet, so dass die tatsächlichen Wertschwankungen einer Geldanlage in Immobilien nicht sachgerecht abgebildet werden. Vor diesem Hintergrund werden offene Immobilienfonds standardmäßig durch einen passenden REIT-Index ersetzt, dessen Bestandteile an einer Börse gehandelt und dementsprechend täglich bewertet werden.

3.3 Bewertung von Anleihen

Bei der Bestimmung des Risikomaßes für Anleihen müssen insbesondere Zinsänderungs- und Ausfallrisiken beachtet werden. Während für ersteres das in Kapitel 2.4.3 vorgestellte Cox-Ingersoll-Ross Modell zum Einsatz kommt, wird letzteres auf Basis der am Markt beobachteten CDS-Spreads hergeleitet. Im einzelnen bedeutet das:

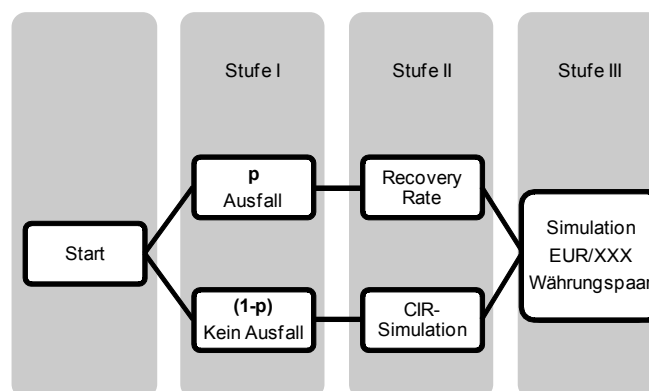


Abbildung 9: Schematische Darstellung des Bewertungsprozesses für Anleihen

- Die Veränderung des maßgeblichen Zinssatzes wird über das CIR-Modell simuliert.
- Die Auswirkung der Zinsbewegung auf die betrachtete Anleihe wird mit Hilfe eines Taylor Polynoms zweiter Ordnung und den Zinssensitivitätskennzahlen **Modified Duration** und **Konvexität** ermittelt.
- Das Ausfallrisiko wird über CDS-Prämien oder im Falle einer fehlenden CDS-Prämie mit Hilfe des in Abschnitt 2.4.4 dargestellten Regressionsmodells geschätzt. Im Falle eines Ausfallereignisses wird die Recovery Rate in Höhe von 40% (Senior) oder 20% (Subordinated) zurückgezahlt.
- Eventuelle Währungsrisiken werden im letzten Schritt über die GARCH-Simulation der Wertentwicklung der zu Grunde liegenden Währung in EUR einbezogen.
- Das Liquiditätsrisiko wird in Form des Bid/Ask-Spreads des entsprechenden Wertpapiers verrechnet. Hierbei wird bei der Bewertung davon ausgegangen, dass das Wertpapier nach dem Betrachtungszeitraum verkauft wird und dabei der Bid/Ask-Spread als Verkaufsabschlag anfällt.

Es ist zu beachten, dass der Bid/Ask-Spread lediglich Aufschluss über die expliziten Liquidationskosten gibt, nicht aber über die versteckten Kosten, die durch eine durch die Verkaufsoffer selbst ausgelöste adverse Marktbewegung entstehen. Die zusätzliche Berücksichtigung der impliziten Transaktionskosten wäre aus fachlicher Sicht zwar wünschenswert, gestaltet sich aber wegen des Fehlens einer breit verfügbaren und vorberechneten Kennzahl als schwierig. Um die impliziten Kosten bzw. das vollständige Liquiditätsrisiko verlässlich zu schätzen, müssten an verschiedenen Börsen mehrere Transaktionen verschiedener Volumina analysiert werden. Dies wiederum würde erhebliche zusätzliche Anforderungen an den Datenbedarf stellen, deren Erfüllung sich nur im Falle von institutionellen Investoren lohnt, die entsprechende Volumina platzieren, da das implizite Kosten berücksichtigende Liquiditätsrisiko stark von der Größe der aufzulösenden Position abhängt. Für Privatanleger dürfte diese Komponente des Liquiditätsrisikos in aller Regel irrelevant sein.

3.4 Bewertung von Einzelaktien

Im Hinblick auf einzelne Aktien wird das Marktrisiko analog zu den Sammelvermögen auf Basis eines GARCH-Modells ermittelt. Darüber hinaus werden weiterhin Ausfall- und Währungsrisiken berücksichtigt. Im einzelnen bedeutet das:

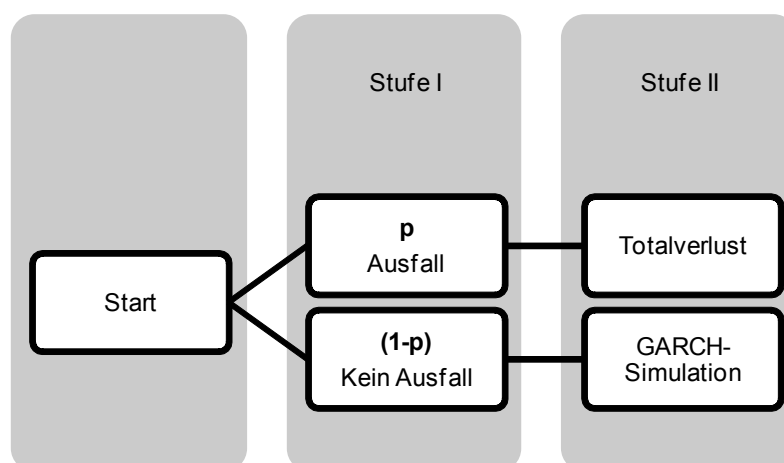


Abbildung 10: Schematische Darstellung des Bewertungsprozesses für Einzelaktien

- Das Marktrisiko wird über ein GARCH-Modell simuliert.

- Ein eventuell vorliegendes Währungsrisiko wird bereits in der Zeitreihe des jeweiligen Wertpapiers berücksichtigt, weil die Modellanpassung auf Basis von Marktdaten in EUR erfolgt, weshalb eine gesonderte Betrachtung in diesem Fall nicht mehr erforderlich ist.
- Das Ausfallrisiko wird über CDS-Prämien oder im Falle einer fehlenden CDS-Prämie über die mittlere geschätzte Ausfallwahrscheinlichkeit der Anleihen des betrachteten Unternehmens approximiert. Falls ein Unternehmen weder mit CDS unterlegt noch mit beobachtbaren Credit Spreads im Sinne von emittierten Anleihen im Markt ist, muss auf einen konservativen Referenzwert wie z. B. das 80%-Quantil aller beobachteten Ausfallwahrscheinlichkeiten zurückgegriffen werden. Im Falle eines Insolvenzereignisses wird ein Totalverlust unterstellt.
- Das Liquiditätsrisiko wird in Form des Bid/Ask-Spreads des entsprechenden Wertpapiers verrechnet. Hierbei wird bei der Bewertung davon ausgegangen, dass das Wertpapier nach dem Betrachtungszeitraum verkauft wird und dabei der Bid/Ask-Spread als Verkaufsabschlag anfällt. Optional könnte bei der Risikobewertung von Einzelaktien auch die Marktkapitalisierung oder der Free Float eines Unternehmens in die Bewertung des Liquiditätsrisikos einbezogen werden, um beispielsweise Penny Stocks innerhalb der Risikoklassen systematisch höher einzuordnen als Blue Chips. So könnten die angesetzten Liquidationskosten in Gestalt des Bid/Ask-Spreads beispielsweise mit einem schrittweise ansteigenden Multiplikator versehen werden, falls Marktkapitalisierung oder Free Float gewisse Grenzwerte unterschreiten.

3.5 Bewertung von Derivaten

Derivate sind wie Aktien mit einem Marktrisiko behaftet, allerdings ist die Partizipation an der Wertentwicklung des Basiswerts in der Regel nicht linear. Die am weitesten verbreiteten Derivate lassen sich jedoch in ein im Sinne des Auszahlungsprofils äquivalentes Portfolio zerlegen, welches sich aus dem Basiswert und / oder verschiedenen Optionen zusammensetzt. Ein Discountzertifikat kann beispielsweise durch den Basiswert ohne Dividendenanspruch als Longposition und eine europäische Calloption als Shortposition abgebildet werden, während sich ein Bonuszertifikat durch den Basiswert ex Dividende und einen Down-and-out Put als Longposition rekonstruieren lässt. Die Bewertung der entsprechenden Optionsscheine kann Hull (2012) entnommen werden.

Darüber hinaus sind Derivate wie auch Anleihen mit einem Emittentenrisiko behaftet. Sie bilden damit die Produktklasse, bei der die Konstruktion von Risikomaßen mit der höchsten Komplexität behaftet ist. Im einzelnen bedeutet das:

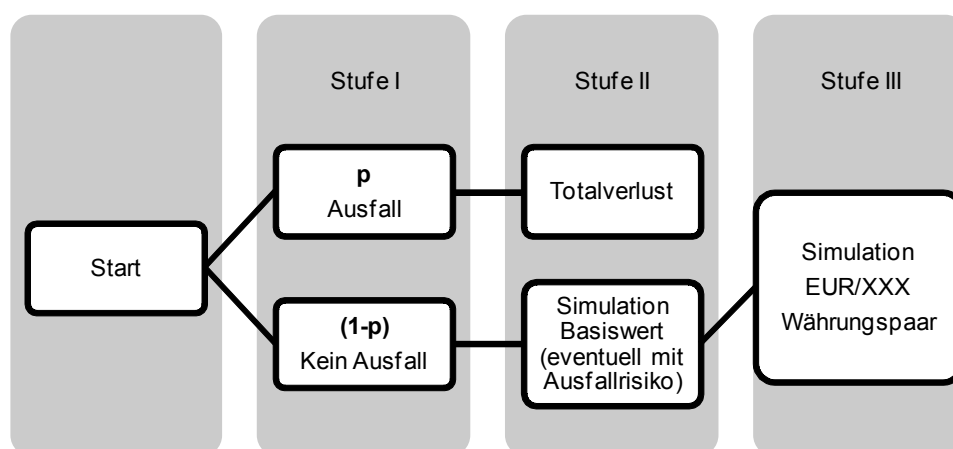


Abbildung 11: Schematische Darstellung des Bewertungsprozesses für Derivate

- Das Marktrisiko wird über ein GARCH-Modell simuliert, wobei es nicht möglich ist, die Wertentwicklung des Derivats direkt zu simulieren. Es ist jedoch möglich, die Wertveränderung des Basiswertes zu simulieren und die korrespondierende Wertentwicklung des Zertifikats auf Basis des funktionalen Zusammenhangs zwischen dem Auszahlungsprofil des Derivats und der Wertveränderung des Basiswerts zu ermitteln.
- Das Ausfallrisiko kann auf Basis der CDS-Prämie geschätzt werden, da alle Emittenten am Markt mit CDS unterlegt sind. Falls eine einzelne Aktie als Basiswert fungiert, ist zu berücksichtigen, dass auch das Ausfallrisiko der zu Grunde liegenden Aktie in das berechnete Risikomaß aufgenommen werden muss. In beiden Fällen, d. h. dem Ausfall des Emittenten und dem Ausfall des Basiswertes, ist in der Regel mit einem Totalverlust zu rechnen. Lediglich Garantiezertifikate sind bei einem Ausfall des Basiswertes nicht von einem Totalverlust betroffen, sondern zahlen zum Laufzeitende den Garantiebetrug aus.
- Das Liquiditätsrisiko wird in Form des Bid/Ask-Spreads des entsprechenden Wertpapiers verrechnet. Hierbei wird bei der Bewertung davon ausgegangen, dass das Wertpapier nach dem Betrachtungszeitraum verkauft wird und hierbei der Bid/Ask-Spread als Verkaufsabschlag anfällt.
- Eventuelle Währungsrisiken werden im letzten Schritt über die GARCH-Simulation der Wertentwicklung der zu Grunde liegenden Währung in EUR einbezogen.

3.6 Exemplarische Bewertung unterschiedlicher Finanzinstrumente

Wendet man die beschriebenen Bewertungsprozesse auf Finanzinstrumente mit unterschiedlichem Ausfallrisiko an, so zeigt die folgende Abbildung für vier exemplarische Finanzinstrumente, zu welchen Ergebnissen man gelangt. Dabei wird insbesondere deutlich, dass sich keineswegs immer eine "Glockenform" ergibt, bei der das Ausfallereignis am linken Verteilungsende die geringste aller Wahrscheinlichkeiten aufweist, sondern dass sich für einige Finanzinstrumente sehr asymmetrische Verteilungen mit teilweise relativ hohen Ausfallwahrscheinlichkeiten ergeben:

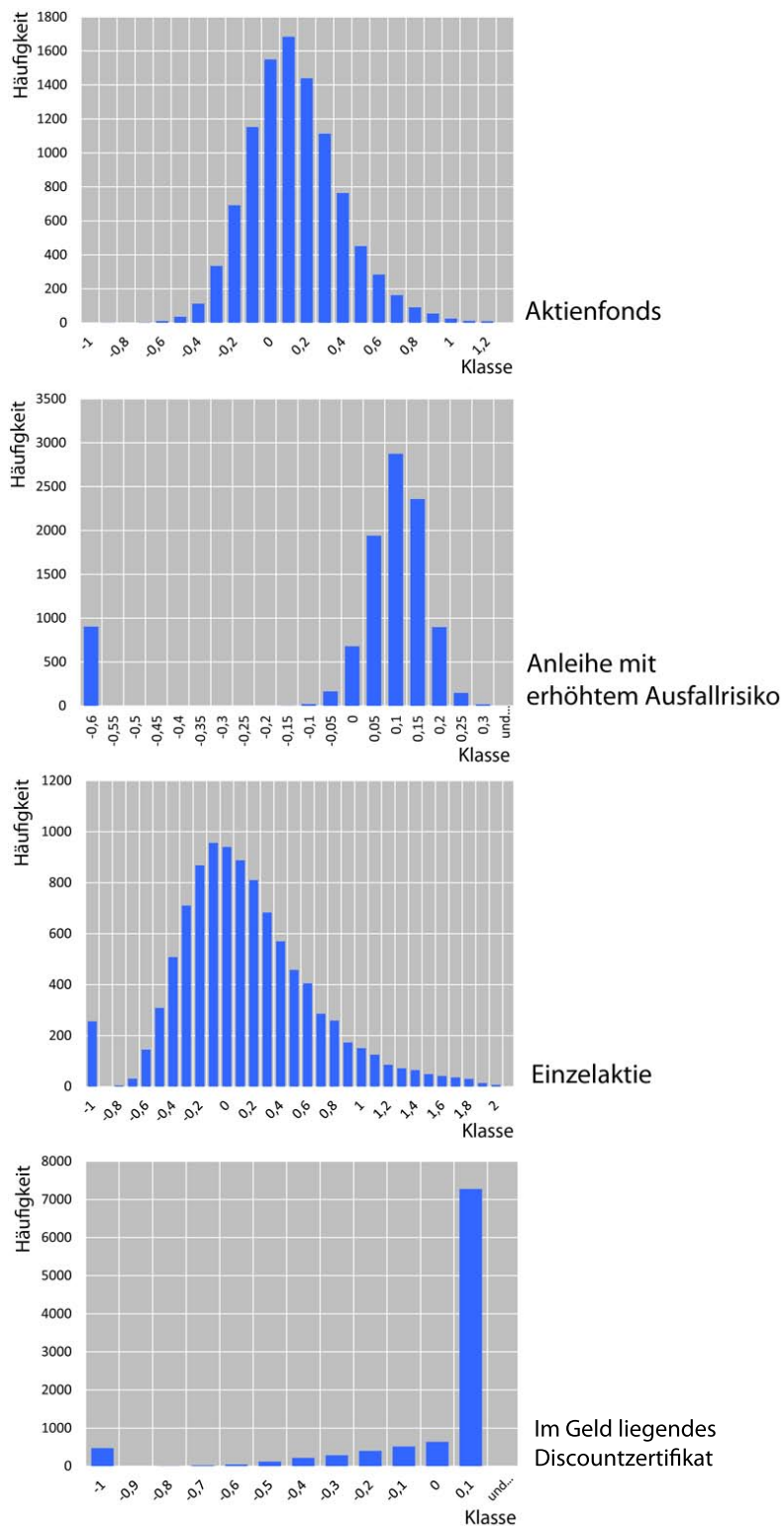


Abbildung 12: Simulierte Wahrscheinlichkeitsverteilungen von Finanzinstrumenten mit unterschiedlichem Ausfallrisiko

4 Risikobewertung bei langer Haltedauer

4.1 Quantitative Merkmale mehrperiodiger Haltedauer

Eine eher selten thematisierte Einschränkung aller bislang vorgestellten Verfahren zur Risikomessung bzw. -modellierung besteht darin, dass sie nur über relativ kurze Zeitspannen eine gute Prognosekraft beanspruchen, d.h. solche, die im Bereich von Tagen, Wochen oder höchstens Monaten liegen. Für viele Belange von institutionellen Investoren ist dieser zeitliche Geltungsanspruch auch durchaus ausreichend, da sie ihre Risiken innerhalb solcher kurzer Zeiträume steuern – entsprechend beziehen sich beispielsweise die Anforderungen an die Eigenkapitalunterlegung der Risikopositionen von Banken innerhalb der Basel-Richtlinien auf einen VaR bei 10-tägiger Haltedauer⁹.

Betrachtet man allerdings die wichtigsten Anlageziele von Privatanlegern und die damit verbundenen Anlagehorizonte, so wird man feststellen, dass diese fast immer deutlich länger sind als Monate, Wochen oder gar Tage. Privatanleger verfolgen ihre Anlageziele in der Regel über Zeithorizonte von mehreren Jahren, und ein Anlageziel von so zentraler Bedeutung wie die Altersvorsorge wird von den meisten Privatanlegern über Zeiträume von weit über 10 Jahren verfolgt. Die daraus resultierende Frage, ob und wie sich die Risikomerkmale von Wertpapieren bei langer Haltedauer ändern und wie ein langfristig orientierter Anleger deshalb sein Portfolio strukturieren sollte, wird in der finanzwissenschaftlichen Literatur merkwürdigerweise nur selten gestellt. Zwei Harvard-Professoren, die sich mit dieser Frage eingehend beschäftigt haben (Campbell and Viceira 2002), vermuten als einen Grund hierfür, dass portfoliotheoretische Forschung bis heute sehr stark von den Grundüberlegen des oben vorgestellten CAPM geprägt ist, bei dem es sich um ein sog. "Einperiodenmodell" handelt. Das heißt, zum Zeitpunkt $t = 0$ wird eine Investition getätigt, die über eine einzige Periode gehalten wird, an deren Ende $t = 1$ beobachtet wird, wie sich der Wert der Investition verändert hat. Würde man die Wertveränderungen der Investition über mehrere Perioden verfolgen – was dem Halten einer Finanzposition über einen langen Zeitraum eher entspricht – so könnten sich daraus bemerkenswerte zusätzliche Erkenntnisse ergeben, wie das folgende einfache Gedankenexperiment zeigt (angelehnt an (Hens and Bachmann 2008):

Nehmen wir an, wir investieren zum Zeitpunkt $t = 0$ einen Betrag von 1.000 Geldeinheiten und es bestehe eine 50%ige Wahrscheinlichkeit, dass dieser Betrag innerhalb einer Periode um 5% steigt, und eine ebenfalls 50%ige Wahrscheinlichkeit, dass der Betrag innerhalb dieser Periode um 3% fällt. Dann wird zum Zeitpunkt $t = 1$ aus der Anfangsinvestition mit 50%iger Wahrscheinlichkeit ein Betrag von 1.050 geworden sein, und mit ebenfalls 50%iger Wahrscheinlichkeit ein Betrag von 970. Betrachten wir nun diese beiden Beträge über eine zweite Periode, von der wir annehmen, dass während ihr die gleichen Wahrscheinlichkeiten für Gewinne und Verluste bestehen. Dann werden wir zum Zeitpunkt $t = 2$ mit 25%iger Wahrscheinlichkeit einen Betrag von 1.103 beobachten, mit ebenfalls 25%iger Wahrscheinlichkeit einen Betrag von 941, und mit 50%iger Wahrscheinlichkeit einen Betrag von 1.019¹⁰. Das heißt, bezogen auf den anfangs investierten Betrag von 1.000 Geldeinheiten werden wir zum Zeitpunkt $t = 2$ nur noch mit 25%iger Wahrscheinlichkeit einen absoluten nominalen Verlust beobachten. Würden wir diese Entwicklung über weitere Perioden beobachten, würde die Wahrscheinlichkeit für einen absoluten nominalen Verlust in der Tendenz noch weiter abnehmen. Diese Abnahme der Wahrscheinlichkeit für einen absoluten nominalen Verlust wäre umso stärker, je größer der positive Erwartungswert wäre, der in den a priori Wahrscheinlichkeiten für Gewinne und Verluste zum Ausdruck kommt.

Solche Überlegen führen zu der Hypothese, dass das Schwankungsrisiko von Wertpapieren mit einer hohen positiven Renditeerwartung (z. B. Aktien) über die Zeit ähnlich wie Anleihen eine Art Mean Reversion hin zu einer positiven langfristigen Durchschnittsrendite aufweisen sollte. Umgekehrt sollte das Schwankungsrisiko von Wertpapieren mit einer niedrigen Renditeerwartung (z. B. Staatsanleihen hoher Bonität) bei langer Haltedauer tendenziell zunehmen.

Campbell und Viceira (2002) haben diese Hypothese anhand von historischen Jahres- und Quartalsrenditen des US-Kapitalmarktes in der Zeit von 1890 bis 1998 für folgende 4 Assetklassen empirisch

⁹Sicherheitswahrscheinlichkeit 99%

¹⁰Nachkommastellen jeweils gerundet

untersucht: (1) rollierte T-Bills (einjährig), (2) rollierte 5-jährige Anleihen, (3) lange Anleihen, die bis zur Endfälligkeit gehalten werden und (4) Aktien. Die folgende Abbildung zeigt die Ergebnisse für Haltedauern bis zu 100 Quartalen, also 25 Jahren.

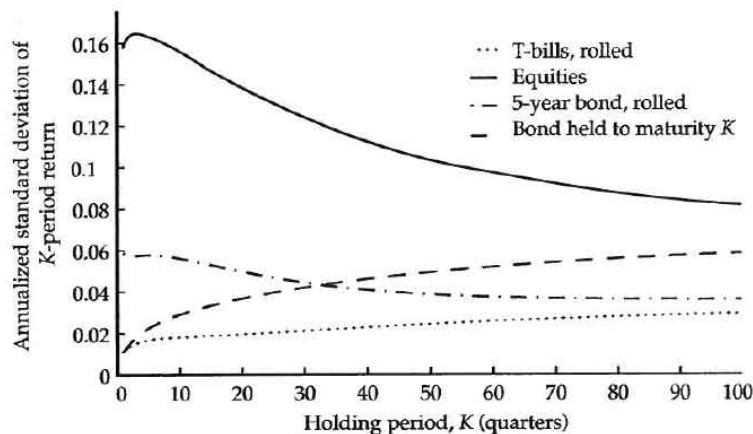


Abbildung 13: Schwankungsrisiko verschiedener Finanzinstrumente in Abhängigkeit von der Haltedauer

Man erkennt, dass die Hypothese weitestgehend bestätigt wird: Das anfangs sehr hohe Schwankungsrisiko von Aktien nimmt bei langer Haltedauer zunehmend ab, während das Schwankungsrisiko von rollierten T-Bills und bis zur Endfälligkeit gehaltenen langen Anleihen tendenziell zunimmt. Die Schwankungsrisiken bei sehr kurzer Haltedauer (nahe der y-Achse) liegen weit auseinander und entsprechen dem, wie man die jeweiligen Assetklassen auch in einem heutigen Risikoklassifikationssystem zuordnen würde. Bei langer Haltedauer sind die zeitraumbezogenen Schwankungsrisiken von Aktien zwar immer noch höher als diejenigen aller Anleihen-Assetklassen, allerdings liegen die Risikoniveaus deutlich enger zusammen.

4.2 Berücksichtigung qualitativer zukünftiger Szenarien

Die bisherigen quantitativen Überlegungen gehen implizit davon aus, dass sich die grundsätzlichen Charakteristika des Schwankungsprofils eines analysierten Finanzprodukts nicht plötzlich verändern, und dass extreme Verlustereignisse rein nach Zufallsprinzip auftreten. In der Realität wird ein professioneller Marktteilnehmer allerdings gerade bei langfristigen Investitionen durchaus qualitativ-ökonomische Hypothesen darüber haben, welche zukünftigen Entwicklungen dazu führen könnten, dass sich die Risikoeigenschaften bestimmter Finanzprodukte grundsätzlich verändern, und in welchen Zeiträumen daher mit außergewöhnlichen Verlustereignissen gerechnet werden muss. Gegen solche außergewöhnlichen Verlustereignisse will sich ein Investor möglichst gezielt wappnen. Einen methodischen Ansatz hierfür bieten ihm sogenannte **Stresstests**.

Stresstests lassen sich in Sensitivitätsanalysen und Szenarioanalysen unterteilen. Sensitivitätsanalysen untersuchen die Wertänderungen einer Risikoposition, die durch eine Änderung eines Marktparameters um eine festgelegte Schrittweite entstehen würden, wie z. B. eine Verschiebung der Zinsstrukturkurve um 50 oder 100 Basispunkte, die Preissteigerung eines Rohstoffes um 20% oder der Kursverfall eines Aktienindex um 30%. Im Unterschied dazu betrachten Szenarioanalysen die Wertänderungen einer Risikoposition bei Eintritt eines Stressszenarios, das in dieser Form entweder historisch bereits aufgetreten ist, oder das in der Zukunft für möglich gehalten wird. Szenarioanalysen betrachten somit ein bestimmtes historisches oder fiktives Muster an Veränderungen von Marktparametern. Im letztgenannten Fall, in dem ein denkbares fiktives Stressszenario betrachtet wird, spricht man auch von "prospektiven" Stress-tests.

Ein Beispiel für einen historischen Stresstest wäre etwa die Untersuchung der Frage, wie eine Risikoposition auf ein erneutes Eintreten derjenigen Marktveränderungen reagieren würde, die beim Platzen der Internetblase oder nach den Terroranschlägen vom 11. September stattgefunden haben. Dadurch, dass das Muster an Marktveränderungen bei historisch realen Ereignissen vollständig bekannt ist, sind historische Stresstests relativ unkompliziert zu berechnen. Sie werfen allerdings inhaltlich die Frage auf, wie realistisch es ist, dass sich bestimmte historische Stressszenarien tatsächlich wiederholen und wie relevant solche Stresstests dementsprechend sind. Deshalb wird im Unterschied dazu bei prospektiven Stresstests der Versuch unternommen, die in naher Zukunft potentiell möglichen Krisenereignisse zu beschreiben und deren Auswirkung auf die untersuchte Risikoposition abzuschätzen. Dabei kann die Konstruktion zukünftiger Szenarien unterschiedlich differenziert vorgenommen werden. Einfachste Varianten des prospektiven Stresstests postulieren beispielsweise festgelegte Maximalverluste für die verschiedenen Anlageklassen eines Depots, und aggregieren diese in gewichteter Form zu einem maximal zu erwartenden Verlust auf Depotebene. Bei der komplexesten Variante des Stresstests wird der Versuch unternommen, die Zusammenhänge zwischen Marktparametern und Anlageklassen bzw. Wertpapieren auf Basis historischer Daten zu analysieren und anschließend unter den Bedingungen eines hypothetischen zukünftigen Stressszenarios zu modellieren. Da es in diesem Fall zu einer kombinierten Anwendung von historischen und prospektiven Stressanalysen kommt, spricht man bei dieser Variante auch von sogenannten "hybriden" Stresstests.

Offenkundig stellt sich bei prospektiven Stresstests die Anforderung, inhaltlich plausible zukünftige Szenarien zu definieren, denen man eine Eintretenswahrscheinlichkeit zubilligt, die oberhalb einer vernachlässigbaren Restwahrscheinlichkeit liegt. Die bisherigen Versuche in der Fachliteratur, solche Stressszenarien für das Finanzsystem zu antizipieren, stammen hauptsächlich von Zentralbanken (Committee on the Global Financial System 2005, Richter 2006). Konkret antizipiert werden dabei beispielsweise ein überraschend starker Konjunkturerinbruch, ein massiver Rückgang des Kreditvergabevolumens oder der Weltgüternachfrage, oder eine massive Auf- bzw. Abwertung bestimmter Währungen. In diesen Szenarien spiegeln sich die makroökonomischen Indikatoren wider, bezüglich derer die Zentralbanken kriseninduzierende Schocks erwarten. Diese Indikatoren lassen sich in enger Anlehnung an entsprechende Modelle der österreichischen Nationalbank (Boss 2002, Kalirai and Scheicher 2002) folgendermaßen kategorisieren:

1. **Konjunktur** (BIP, Industrieproduktion, Produktionslücke)
2. **Preisstabilität** (Verbraucher-, Energie-, Rohstoffe-, Immobilienpreise, Mietniveau, Geldmengenwachstum)
3. **Staat** (Verschuldung, Staatsquote, Bevölkerungsentwicklung, Politische Stabilität)
4. **Banken / Unternehmen** (Kreditvergabe, Unternehmensgewinne, Industrieproduktion, Investitionsausgaben, Bruttoanlageinvestitionen, Insolvenzen, Geschäftsklima)
5. **Private Haushalte** (Erwerbseinkommen, Konsumausgaben, Kfz-Neuzulassungen, Arbeitslosenrate, Sparquote, Verschuldung, Lohnniveau, Verbrauchervertrauen)
6. **Finanzmarkt** (Leitzins, kurz-/langfristige Nominal- und Realzinsen, Zinsstrukturkurve, Hypothekenschulden, Kursniveau Aktienmärkte, Volatilitätsniveau Aktienmärkte)
7. **Externe Faktoren** (Welthandel, Wechselkurse, Exporttätigkeit, ökologische Bedingungen)

Auf diesen Ebenen definieren die Zentralbanken zur Durchführung prospektiver Stresstests hypothetische Szenarien, die üblicherweise einen Antizipationshorizont von ein bis maximal drei Jahren besitzen. Darin spiegelt sich noch einmal der oben bereits angesprochene Sachverhalt wider, dass institutionelle Risikosteuerung eine durchschnittliche Haltedauer von Finanzinstrumenten impliziert, die deutlich kürzer ist als der Anlagehorizont der meisten Privatanleger. Ein Privatanleger, der beispielsweise seine Geldanlagen für die Altersvorsorge aus einer Stressperspektive betrachten möchte, muss sich daher eher die Frage stellen, welche längerfristigen Risikofaktoren für den Kapitalmarkt bestehen und wie sich diese auf seine Kapitalanlagen auswirken können. Studien, die sich mit solchen Fragestellungen

beschäftigen, existieren durchaus. Sie werden von großen geldpolitischen Institutionen wie den Zentralbanken, dem Internationalen Währungsfonds, der Bank für Internationalen Zahlungsausgleich, der Europäischen Kommission etc. regelmäßig erstellt und veröffentlicht. Stressszenarien, die dabei bereits thematisiert wurden, sind beispielsweise die Zahlungsunfähigkeit eines Staates im Euroraum (European Commission Directorate-General for Economic and Financial Affairs 2002), ein dauerhafter Nachfragerückgang nach Gütern aller Art als Folge einer ungünstigen demographischen Entwicklung (z. B. Just 2009), oder anhaltend hohe finanzielle Belastungen für die Bewältigung von ökologischen Veränderungen (Spalding 2010). Es handelt sich dabei um solche Szenarien, die zu einem bestimmten Zeitpunkt in Umrissen bereits erkennbar sind, zur Entfaltung ihres vollen Krisenpotenzials aber eine längere Entwicklungszeit benötigen, dann allerdings viele der oben genannten Ebenen gleichzeitig erfassen und sich insofern in übergeordnete Probleme verwandeln würden.

Obwohl gerade solche Studien aus der Sicht eines längerfristig orientierten Investors letztlich am relevantesten sind, finden sie bislang keine Berücksichtigung bei der Durchführung von Stresstests, was mehrere Gründe haben dürfte: Einer ist sicher der oben bereits angesprochene Sachverhalt, dass Stresstests primär im Rahmen der institutionellen Bankensteuerung und -regulation eingesetzt werden und dabei üblicherweise sehr kurze Haltedauern von Risikopositionen betrachtet werden. Ein weiterer Grund dürfte sein, dass jede Verlängerung des Antizipationshorizonts sowohl den rechnerischen als auch den intellektuellen Aufwand bei der Definition zukünftiger Stressszenarien erhöht, da sich das Spektrum möglicher Entwicklungen mit zunehmendem Antizipationshorizont immer weiter aufspreizt. Deshalb wäre es sicherlich auch vermessen, ein Modell aller denkbaren langfristigen Stressszenarien des Kapitalmarkts entwickeln zu wollen. Wir denken allerdings, dass es möglich ist, innerhalb eines klar abgegrenzten Geltungsbereichs zumindest solche langfristig denkbaren Szenarien zu definieren, die dann eintreten, wenn sich bestimmte Entwicklungen fortsetzen, die heute bereits in Ansätzen erkennbar sind. Dabei grenzen wir den Geltungsbereich folgendermaßen ein:

1. **Beschränkung auf kumulative Krisen:** In der Krisenforschung wird zwischen sogenannten "abrupten" und "kumulativen" Krisen unterschieden (Sigrist et al. 2010). Abrupte Krisen zeichnen sich dadurch aus, dass sie jederzeit auftreten können, gleichzeitig aber ihrem Wesen nach nicht vorhersehbar sind. Beispiele für abrupte Krisen wären etwa Terroranschläge, Naturkatastrophen oder Seuchen. Ereignisse dieser Art bilden ein permanent vorhandenes Hintergrundrisiko, das bezüglich des Zeitpunkts seines Auftretens und seiner Ausprägung nicht näher bestimmt werden kann. Kumulative Krisen zeichnen sich demgegenüber dadurch aus, dass sie sich sukzessive über die Zeit aufbauen, sich dabei verstärken und nach Überschreitung eines kritischen Punktes schlagartig ihr akkumuliertes Potenzial entladen. Kumulative Krisen verhalten sich somit wie große Wellen, die sich über dem Meer langsam aufbauen und beim Auflaufen auf ein Riff brechen. Sie besitzen somit eine gewisse zeitabhängige Regelmäßigkeit, durch die die Bedingungen ihres Auftretens näher bestimmt werden können. Dies wiederum macht sie insgesamt besser vorhersagbar und somit für prospektive Stresstest besser verwendbar.
2. **Marktrelevanz:** Prospektive Stresstests können sich nur mit solchen Krisen beschäftigen, deren Auftreten eine unmittelbare Auswirkung auf die Preise von Finanzinstrumenten hätte. Es werden somit keine Phänomene betrachtet, die ebenfalls als "Krise" bezeichnet werden, allerdings keine oder eine zumindest völlig unklare Wirkung auf die Preise von Finanzinstrumenten besitzen.
3. **Reichweite:** Prospektive Stresstests sollten sich auf solche Krisen konzentrieren, die einen erheblichen Teil des globalen Kapitalmarktes betreffen und dort per saldo substantielle Vermögenswerte vernichten würden. Damit trägt man dem Umstand Rechnung, dass in bestimmten Teilsegmenten des globalen Kapitalmarktes (z.B. in einzelnen Unternehmen, Ländern oder Branchen) in gewisser Weise ständig kleinere Krisen stattfinden, die in anderen Teilsegmenten allerdings zu Vermögenswachstum führen.
4. **Anlegerrelevanz:** Prospektive Stresstests sollten den Risiko- und den Verlustbegriff aus der Nutzenperspektive eines Kapitalanlegers verwenden und somit in einer Weise, wie es die sogenannte "Prospect Theory" von Kahneman and Tversky (2010) nahelegt. Somit sollte der Begriff des "Risikos" nicht im Sinne der (positiven wie negativen) Kursschwankungsbreite einer Kapitalanlage

verstanden werden, sondern im Sinne der Gefahr, mit einem Anlageprodukt am Ende einer geplanten Haltedauer einen Verlust zu erleiden und somit das Ziel der Kapitalanlage zu verfehlen. Ein "Verlust" wird dementsprechend nicht an der nominellen Höhe des investierten Betrages bemessen, sondern an dessen Kaufkraft, also inflationsbereinigt (= real).

5. **Eindeutige Währungsperspektive:** Prospektive Stresstests sollten stets klar machen, aus welcher Währungsperspektive sie eine Krise definieren, da sich manche Entwicklungen nur in bestimmten geographischen Regionen und deren Währung als Krisen manifestieren würden.
6. **Krisenüberlagerung:** In der Realität werden große Krisen nie isoliert auftreten, sondern an ihren Rändern wiederum krisenhafte Seiten- und Folgephänomene bewirken. So hat beispielsweise die Bankenkrise des Jahres 2008 in der Folge auch eine Staatsschuldenkrise ausgelöst. Insofern sollten "Hauptkrisen" unterschieden werden, die auf verschiedenen Ebenen entstehen können. Dass es beim tatsächlichen Ausbruch einer Hauptkrise dann faktisch immer zu einer Überlagerung verschiedener Krisenarten käme, soll damit nicht angezweifelt werden.
7. **Ausschluss eines totalen Zusammenbruchs:** Prospektive Stresstests können sich auf die Betrachtung von solchen Krisen beschränken, die im Rahmen des uns heute bekannten globalen Wirtschafts- und Finanzsystems potenziell überwindbar wären und nicht zu dessen vollständigem Zusammenbruch führen würden, da in diesem Fall - dem logisch ultimativen "Worst Case" - die angestellten Überlegungen ohnehin obsolet würden.

Studiert man innerhalb eines solchen Rahmens die hypothetischen Krisenszenarien, die von großen geldpolitischen Institutionen diskutiert werden, so lässt sich daraus eine Art "Risikoatlas" möglicher kurz-, mittel- und langfristiger Entwicklungen erstellen. Im Rahmen einer Studie in Zusammenarbeit mit der Deutschen Schutzvereinigung für Wertpapierbesitz e.V. (DSW) haben wir im Jahr 2010 einen solchen Risikoatlas für den Euroraum konkret dargestellt (Beck and Layes 2010, S. 29). Er beinhaltet ein Basisszenario und 4 Stressszenarien. Das Basisszenario entspricht einer weitgehend schockfreien Entwicklung, innerhalb derer sich die Verlustpotenziale in einem Rahmen bewegen, der durch eine Normalverteilung annähernd gut beschrieben werden kann. Demgegenüber stellen die Stressszenarien krisenhafte Entwicklungen dar. Mit Hilfe qualitativ-ökonomischer Überlegungen wurde dann jedes dieser Szenarien mit einer Eintrittswahrscheinlichkeit und jede Assetklasse mit einem "Verlустаufschlag" belegt, d.h. mit einem Verlust, der bei dieser Assetklasse beim Eintritt des jeweiligen Szenarios zusätzlich zu dem Verlust zu erwarten ist, der im Rahmen einer klassischen VaR-Berechnung auf Basis einer Monte-Carlo-Simulation zu erwarten ist.

Ein solches Vorgehen eröffnet die Möglichkeit, ein hybrides Risikomaß zu bestimmen, d.h. ein Risikomaß, das einerseits die historischen Renditen eines Finanzprodukts einbezieht, gleichzeitig aber auch berücksichtigt, dass dieses Finanzprodukt möglicherweise sehr lange gehalten wird und in dieser Zeit von verschiedenen Stressszenarien betroffen sein kann, die heute prognostiziert werden. Wenn dabei außerdem ein Verlust nicht an der nominellen Höhe, sondern an der Kaufkraft des investierten Betrages bemessen wird, so trifft ein solches hybrides Risikomaß klare Aussagen zum langfristigen monetären Nutzen eines Anlageprodukts und stellt somit für den langfristig orientierten Anleger eine bessere Entscheidungshilfe dar als ein einzelnes klassisches Risikomaß.

Literatur

- Artzner, P.; Delbaen, F.; Eber, J. M., and Heath, D. Thinking coherently. *RISK*, 10(11):68 – 71, 1997.
- Beck, A. and Layes, G. *Mit Sicherheit zu wenig: Das Dilemma der privaten Altersvorsorge. Studie zum Nutzen risikobehafteter Finanzprodukte und Ableitung eines zeitsensitiven Risikomaßes*. Institut für Vermögensaufbau, 2010.
- Bollerslev, T. Generalized autoregressive conditional heteroskedasticity. *Journal of Econometrics*, 31: 307 – 327, 1986.
- Boss, M. *Finanzmarktstabilitätsbericht 4, Ein makroökonomisches Kreditrisikomodell zur Durchführung von Krisentests für das österreichische Kreditportfolio*. Österreichische Nationalbank, 2002.
- Campbell, J. Y. and Viceira, L. M. *Strategic Asset Allocation, Portfolio choice for Long-Term Investors*. Oxford University Press, 2002.
- Chen, N. F.; Roll, R., and Ross, S. A. Economic forces and the stock market. *Journal of Business*, 59: 383–403, 1986.
- Committee on the Global Financial System, . *Stress testing at major financial institutions: survey results and practice*. BIS, 2005.
- Cox, J. C.; Ingersoll, J. E., and Ross, S. A. A theory of the term structure of interest rates. *Econometrica*, 53(2):385 – 408, 1985.
- European Commission Directorate-General for Economic and Financial Affairs, . *Sustainability Report 2009*. European Commission, 2002.
- Hackethal, A.; Inderst, R.; Meyer, S., and Rochow, T. *Messung des Kundennutzens der Anlageberatung. Wissenschaftliche Studie im Auftrag des Bundesministeriums für Ernährung, Landwirtschaft und Verbraucherschutz (BMELV)*. 2011.
- Heath, D.; Jarrow, R., and Morton, A. Bond pricing and the term structure of interest rates: A new methodology for contingent claims valuation. *Econometrica*, 60, 1992.
- Hens, T. and Bachmann, K. *Behavioural finance for private banking*. Wiley Finance, 2008.
- Ho, T. S. Y. and Lee, S. B. Term structure movements and pricing interest rate contingent claims. *Journal of Finance*, 41, 1986.
- Hull, J. C. *Options, Futures and other Derivatives*. Pearson Education Limited, 8. edition, 2012.
- Hull, J. C. and White, A. Pricing interest-rate derivative securities. *The Review of Financial Studies*, 3, 1990.
- Jorion, P. *Financial Risk Manager Handbook*. John Wiley and Sons, 6. edition, 2011.
- Just, T. *Demografie und Immobilien*. Oldenbourg, 2009.
- Kahneman, D. and Tversky, A. Prospect theory: An analysis of decision under risk. *Econometrica*, 47(2): 263–291, 2010.
- Kalirai, H. and Scheicher, M. *Finanzmarktstabilitätsbericht 3: Makroökonomische Krisentests: Erste Ergebnisse für Österreich*. Österreichische Nationalbank, 2002.
- Lazard Asset Management, . *Vom unvermeidlichen Scheitern des Risikomanagements*. 2009.
- Markowitz, H. M. Portfolio selection. *Journal of Finance* 7, pages 71–91, 1952.

Richter, J. *Stresstests in Banken. Von Basel II bis ICAAP, Stresstests in der Fachliteratur.* Klauck & Stegmann, 2006.

Sharpe, W. Capital asset prices – a theory of market equilibrium under conditions of risk. *Journal of Finance* XIX (3), 1964.

Sigrist, T.; Sweys, M., and Vögeli, R. *Krisentheorien.* Institut für Schweizerisches Bankenwesen der Universität Zürich, 2010.

Spalding, K. S. *Investors Analyze Climate Risks and Opportunities: A Survey of Asset Managers' Practices.* Ceres, 2010.

Taleb, N. N. *The Black Swan: The Impact of the Highly Improbable.* Penguin, 2008.