

MASTERARBEIT

Titel der Masterarbeit

„Wieviel Trend steckt im Zufall?“

Verfasser

Florian Göschl, Bakk. rer. soc. oec.

angestrebter akademischer Grad

Master of Science (MSc)

Wien, 2015

Studienkennzahl lt. Studienblatt:
Studienrichtung lt. Studienblatt:
Betreuer:
Mitbetreuer:

A 066 915
Masterstudium Betriebswirtschaft
o. Univ.-Prof. Dipl.-Math. Dr. Jörg Finsinger
Mag. Dr. Stephan Unger

1 Inhaltsverzeichnis

1	Inhaltsverzeichnis.....	3
2	Einführung.....	4
3	Die Theorie der effizienten Märkte.....	5
3.1	Kritikpunkte der Market-Efficiency-Hypothese	7
4	Die „Technische Analyse“ der Finanzmärkte	9
4.1	Kritikpunkte der „Technischen Analyse“	11
5	Die Trendstrukturanalyse der simulierten Daten	14
5.1	Der Algorithmus zur Untersuchung von Trendstrukturen in simulierten Daten.....	17
5.1.1	Die Erzeugung von Trendgeraden.....	18
5.1.2	Die Initiierung neuer Trends bei Durchbrechen von Preisbewegungen	20
5.1.3	Die Initiierung neuer Trends ohne Durchbrechen von Preisbewegungen.....	21
5.1.4	Die Erzeugung sämtlicher Trendgeraden innerhalb einer geometrischen Brownschen Bewegung.....	23
6	Statistische Auswertung der simulierten Daten	24
7	Die Trendstrukturanalyse des EURO STOXX 50.....	29
8	Statistische Auswertung der EURO STOXX 50 Daten	31
9	Simulierte Daten vs. Historische Daten	33
10	Zusammenfassung.....	38
11	Quellenverzeichnis	40
11.1	Bücher	40
11.2	Zeitschriften	40
11.3	Online-Quellen.....	41
11.4	Dissertationen.....	41
12	Anhang	42
12.1	R-Code	42
12.2	Abstract (Deutsch).....	46
12.3	Abstract (English)	46
12.4	Curriculum Vitae.....	47
12.5	Eidesstattliche Erklärung.....	48

2 Einführung

Schon seit dem 18. Jahrhundert wird versucht Zufallsprozesse in Modellen zu beschreiben. Bekannte Pioniere auf dem Gebiet waren unter anderem der Biologe Robert Brown um 1830 bzw. der Mathematiker Louis Bachelier. Letzterer beschrieb bereits um 1900 Kursschwankungen als Zufallsprozesse und sah sich damals schon einigen Kritikern in der Wissenschaft gegenüber.

Der Versuch Aktienkursbewegungen zu beschreiben und vorherzusagen war seit dem immer ein großes Thema bei Experten der Mathematik und Finanzwirtschaft. Deshalb widmet sich die vorliegende Arbeit diesem Thema und behandelt zunächst die zwei großen Denkschulen, die diesbezüglich ganz verschiedene Ansätze haben. Daher wird zunächst die mathematisch-statistische Aktienanalyse besprochen, welche vor allem in der Finanztheorie angesiedelt ist und sowohl klassische Ansätze, wie die Market-Efficiency-Hypothese beinhaltet, als auch modernere Theorien des Chaos untersucht. Dem gegenüber steht die „Technische Analyse“ der Finanzmärkte. Diese Ansätze gehen auf Charles Dow zurück und beschreiben Aktienkurse weniger mathematisch, sondern viel mehr in Form von Mustern, Trends, Zyklen, etc. .

Die Konkurrenz beider Denkschulen ist bis heute ungebrochen, denn bis dato konnte noch kein Ansatz weder bestätigt noch widerlegt werden. Deshalb wird in dieser Arbeit versucht, mit Hilfe eines Algorithmus, Trends mit verschiedener Länge in simulierten und historischen Aktienkursen zu untersuchen. Würde man dabei die Basis einer Gesetzmäßigkeit finden, könnte man in weiterer Folge im Sinne der „Technischen Analyse“ statistische Vorhersagen bezüglich der zu erwartenden Trends geben. Zusätzlich wird die durchschnittliche Trendstruktur der simulierten Daten mit der durchschnittlichen Trendstruktur der historischen Daten verglichen um Rückschlüsse auf den Einfluss vom Zufall auf Aktienkurse zu ziehen.

3 Die Theorie der effizienten Märkte

Einer der wichtigsten Grundsätze der derzeit allgemein akzeptierten ökonomischen Finanzmarkttheorie ist die Theorie der „Effizienten Märkte“ bzw. Market-Efficiency-Hypothesis (MEH). Wie in anderen Disziplinen, die sich mit der Analyse bzw. Beschreibung der Finanzmärkte auseinandersetzen, wird auch hier versucht das Verhalten der Marktteilnehmer bzw. der Renditen und Marktpreise zu verstehen und zu beschreiben.¹

Dabei werden vor allem die Annahmen des sogenannten „Urnenmodells“ aufgegriffen, wonach die Renditen einzelner Jahre unabhängige Ziehungen sind, was auch bedeutet, dass sie unkorreliert sind. Dies war allerdings vielen Theoretikern und Praktikern um 1960 ein Dorn im Auge, da man ökonomisch keine Erklärung für die Zufälligkeit finden konnte. Etwa zur selben Zeit wurde die Widersprüchlichkeit zwischen zufälligen Renditen und der „Rationalität der Investoren“ scheinbar aufgelöst. Davor bestand die Frage, warum es relativ lange dauern sollte, bis uninformierte Marktteilnehmer neue Informationen verarbeiten bzw. den informierten Marktteilnehmern Glauben schenken sollten. Allerdings kam man mit der Einsicht, dass nur schnell agierende Marktteilnehmer Geld verdienen können, zu dem Schluss, dass Investoren unmittelbar auf Informationen reagieren.²

Darüber hinaus versuchen sie auch zu antizipieren, was bedeutet, dass sie nicht nur die konkreten Informationen aufnehmen, sondern sich auch überlegen was diese Informationen im weiteren Verlauf bewirken könnten. An der Börse sind dann besonders jene Personen aktiv, die diese Informationen in neue Werte für Aktien und Bonds übersetzen können. Dementsprechend werden dann vor allem in großen Volumen entsprechende Positionen auf- und abgebaut. An einem bestimmten Punkt hat dann der Kurs durch die Informationsauswertung und über Angebot und Nachfrage das korrekte Wertniveau erreicht. Die ganze Informationsverarbeitung passiert sehr schnell und gewissenhaft, weil bei den Investitionen viel Geld auf dem Spiel steht. Das bedeutet, die Marktteilnehmer sind rational und so schnell, dass etwa Preise von Wertpapieren so gut wie unmittelbar nach Eintreffen einer neuen Nachricht dem korrekten Niveau entspricht. Natürlich müssen die Nachrichten unerwartet sein. Denn so finden sich im Finanzmarkt stets die „korrekten“ Werte welche alle Informationen und Erwartungen enthalten.³

An diesem Punkt stellt sich noch die Frage welche Arten von Informationen in den Marktwerten enthalten sind. Nach Fama (1970) kristallisieren sich dabei drei Arten der Informationseffizienz ab. Danach kann ein Markt schwach, semi-stark oder stark informationseffizient sein. Bei der

¹ Vgl. Thoma (1997), S. 6

² Vgl. Spremann (2006), S. 154

³ Vgl. Spremann (2006), S. 154f.

schwachen Informationseffizienz sind bloß historische Zeitreihen und Kurse verfügbar, während bei der semi-starken bzw. starken Informationseffizienz noch sämtliche veröffentlichte Informationen bzw. noch zusätzlich Insiderwissen dazukommt.⁴

Die Annahmen der Market-Efficiency-Hypothese, dass etwa Aktienkurse völlig zufällig entstehen, scheinen in der Praxis auch zunächst zu stimmen. Oftmals reagieren Kurse beispielsweise gar nicht auf positive Nachrichten, weil positive Ereignisse bereits erwartet und somit eingepreist wurden. Aber auch bei überraschenden Nachrichten kann man oft sehr schnell Reaktionen erkennen. Bevor der durchschnittliche Marktteilnehmer überhaupt zum Zug kommt, hat sich der Kurs bereits geändert, was also bedeutet, dass der Markt effizient reagiert.⁵

Die Konsequenzen der Theorie der effizienten Märkte sind ein Hauptthema in dieser Arbeit. Werden nämlich alle Informationen umgehend und in vollem Umfang in den Marktwerten verarbeitet, lassen sich Informationen aus der Vergangenheit nicht für Gewinne in der Zukunft nutzen. Bewegen sich Aktien also völlig zufällig, kann man keine Gesetzmäßigkeiten dieser Bewegung finden und auch nicht von der Vergangenheit auf die Zukunft schließen.⁶

Modelliert wird der Zufall der effizienten Märkte unter anderem durch den sogenannten „Random Walk“. Dieser kann nach Thoma (1997) auch recht einfach durch ein Würfelexperiment beschrieben werden. Dabei liegt etwa der Startwert einer Aktie beispielweise bei 100. Würfelt man nun die Zahlen 1, 2 oder 3, erhöht sich der Wert der Aktie um eine Einheit. Trifft man allerdings die Zahlen 4, 5, oder 6, so wird der Wert der Aktie um eine Einheit gemindert. Ob der Preis der Aktie zulegt oder abnimmt hängt also nur vom Zufall ab. Würfelt man nun öfter, ergibt sich ein Kurs im Koordinatensystem, welcher einem echten Aktienchart nicht unähnlich ist. Interessant ist allerdings was passiert, wenn man so einen Chart in der Art und Weise viel öfter erzeugt und die Unterschiede zwischen den jeweiligen Startwerten und den Endwerten des Aktienkurses analysiert. Die Abweichungen sind stark normalverteilt, was bedeutet, dass der Kurs praktisch immer um den Startwert herum pendelt, weil die Wahrscheinlichkeit für eine Aufwärtsbewegung stets gleich groß ist, wie für eine Abwärtsbewegung. Der wichtigste Schluss der sich aus dieser Tatsache ergibt ist, dass also keine (längerfristigen) Trends innerhalb eines Kurses zu erwarten sind und dass Informationen aus der Vergangenheit keinerlei Einfluss auf die Zukunft haben.⁷ Sobald sich die Kurse also angepasst haben, laufen sie im Allgemeinen seitwärts.⁸ Die Annahme der Effizienz der

⁴ Vgl. Spremann (2006), S. 158

⁵ Vgl. Thoma (1997), S. 7

⁶ Vgl. Thoma (1997), S. 7f.

⁷ Vgl. Thoma (1997), S. 12f.

⁸ Vgl. Thoma (1997), S. 45

Finanzmärkte wurde in der Vergangenheit zwar empirisch überprüft und einigermaßen bestätigt, aber nicht eindeutig bewiesen.

3.1 Kritikpunkte der Market-Efficiency-Hypothesis

Wie im Folgenden genauer beschrieben wird, gibt es allerdings Widerstände gegen die Market-Efficiency-Hypothesis. Zunächst wird direkt die Informationseffizienz relativiert. Aus der Praxis weiß man, dass Informationen eben nicht immer sofort verarbeitet werden bzw. erst zu späteren Zeitpunkten zu Handlungen führen. Somit können konkrete Entscheidungen erst Minuten, Stunden, Tagen oder Monaten getroffen werden. Manchmal führen neue Informationen auch zu gar keinen Handlungen, sondern ändern nur die Grundeinstellung eines Marktteilnehmers. Spätestens hier wirkt die Vergangenheit sehr wohl auf die Zukunft und die völlige Unabhängigkeit der Ereignisse ist nicht mehr gewährleistet. Solche Situationen beschreibt der Random Walk nicht mehr.⁹

Abgesehen davon zweifelte bereits Albert Einstein an der sehr einfachen Sicht auf den Zufall. Er misstraute dem leichtfertigen Umgang mit Random Walks und Normalverteilungen in diesem Zusammenhang, konnte aber Zeit seines Lebens keine Fortschritte auf dem Gebiet machen. Dabei ist es dennoch einleuchtend, dass gerade in stark vernetzten Systemen wie Gesellschaft oder Finanzmärkten die Unterstellung der völligen Unabhängigkeit durchaus gewagt scheint. Ganz besonders Technische Analysten bzw. Charttechniker sind entschiedene Gegner der Theorie der effizienten Märkte. Bedenkt man nochmals, dass die MEH von keinerlei Trends in den Finanzmärkten ausgeht, jedoch Technische Analysten genau solche vorherzusagen versuchen, um etwaige Kauf- oder Verkaufsempfehlungen zu geben, scheint dies einleuchtend.¹⁰

Bevor allerdings die „Technische Analyse“ der Finanzmärkte diskutiert wird, soll zunächst kurz auf die Theorie der durchschnittlichen Schwankungsbreiten eingegangen werden. Oftmals interessiert man sich bei der Preisentwicklung von Aktien nämlich für die Schwankungsbreiten bzw. Volatilität und wie sie sich über die Zeit entwickelt. In den Untersuchungen von Aktienkursen wird statt der Volatilität häufig noch ein anderes Streuungsmaß verwendet. Die sogenannte „Rescaled Range“ (R/S) ist der Varianz grundsätzlich sehr ähnlich, in der Berechnung der Abweichungen vom Mittelwert werden allerdings alle Zahlen einfach aufsummiert, woraus eine Folge von Zahlen entsteht. Diese Folge hat einen maximalen bzw.

⁹ Vgl. Thoma (1997), S. 46f.

¹⁰ Vgl. Thoma (1997), S. 47

minimalen Wert dessen Differenz die „Range“ ist. „Reskaliert“ bedeutet weiterhin, dass die zuvor berechnete Schwankungsbreite noch durch die gewöhnliche Standardabweichung dividiert wird, was den Vergleich verschiedener Aktien verbessert.¹¹

Das Interessante ist nun aber vor allem, wie verschiedene Schwankungshorizonte miteinander zusammenhängen. Die Frage, die hier zugrunde liegt ist, ob beispielsweise Wochenschwankungen fünfmal höher sind als Tagesschwankungen. Das wäre zwar ein logischer Gedanke, er ist aber dennoch falsch. Konkret würde die durchschnittliche Schwankungsbreite nicht um fünf sondern um die Quadratwurzel aus fünf anwachsen. Folgende Formel zeigt den Zusammenhang zwischen der Rescaled Range und Zeit, wobei n die Anzahl der Tage repräsentiert und c eine Konstante darstellt, die hier gleich 1 ist:

$$(R/S)_n = c\sqrt{n} = cn^{\frac{1}{2}}$$

Der Fokus liegt nun im weiteren Verlauf vor allem auf dem Exponenten (1/2 bzw. 0,5). Gemäß der Portfoliotheorie (d.h. bei reinem Zufall) würde ein Exponent von 0,5 zutreffen, allerdings fand Harold Hurst, der Schwankungen von Wasserständen am Nil untersuchte, heraus, dass die obige Formel so nicht der Realität entsprach. Er modifizierte also den Exponenten („Hurst Exponent“ bzw. „H“), wobei er wusste, dass dieser alle Zahlen zwischen 0 und 1 einnehmen kann. Hurst fand heraus, dass H mit einem Wert von 0,7 die Realität am besten beschrieb. Diese konkrete Formel beschrieb also recht gut den Zusammenhang zwischen einem bestimmten Zeitintervall und der jeweiligen durchschnittlichen Schwankungsbreite des Nil.¹²

Die Frage ist nun, ob die Erkenntnisse von Hurst auch an den Finanzmärkten Gültigkeit besitzen. Hätten Aktienkurse, gemäß der totalen Unabhängigkeit, dieselben Eigenschaften wie ein Random Walk und würde man die Schwankungsbreiten in ein sogenanntes „Doppelloarithmisches Diagramm“ eintragen, würde gemäß dem entsprechenden Hurst Exponent von 0,5 auch eine Gerade mit der Steigung von 45% entstehen. Nach Thoma (1997) kann aber nach der R/S-Analyse eines internationalen Aktienindex von 1959 bis 1990 ein anderes Ergebnis festgestellt werden. Die Gerade im Diagramm folgt einer Steigung von 0,7, was bedeutet, dass auch H gleich 0,7 ist. Die Analyse anderer Aktien und Indizes bringt ähnliche Ergebnisse, wobei die Exponenten stets über 0,5 liegen.¹³ Bemerkenswert ist dabei, dass Hurst-Exponenten welche Zahlen zwischen 0,5 und 1 annehmen, auf tendenzielle Aufwärts- bzw.

¹¹ Vgl. Thoma (1997), S. 27

¹² Vgl. Thoma (1997), S. 31ff.

¹³ Vgl. Thoma (1997), S. 47ff.

Abwärtsbewegungen schließen lassen. Je weiter sich die Werte von 0,5 wegbewegen und in Richtung 1 wandern, desto stärker ausgeprägt ist der Trend und desto „glatter“ verläuft die Bewegung.¹⁴ Je kleiner H ist, desto gezackter fällt die Preisbewegung aus. Dabei werden die Trends sehr früh abgeblockt und bewegen sich stark um den Startwert. Jedenfalls sprechen Trends, welche in der Praxis mit H zwischen 0,5 und 1 liegen, grundsätzlich gegen die reine Unabhängigkeit von Ereignissen und stehen damit ganz klar im Widerspruch zur Market-Efficiency-Hypothese.¹⁵ Zu einem sehr ähnlichen Schluss kommt auch Hofmann (1973). Er konnte Trends, d.h. Preisbewegungen welche sich (statistisch) stärker vom Startwert wegbewegen, signifikant nachweisen und lehnte somit die „Random Walk“ Hypothese ab. Je nach gefordertem Realitätsgrad ist diese als Modellvorstellung aber dennoch brauchbar.¹⁶

4 Die „Technische Analyse“ der Finanzmärkte

Wie bereits oben erläutert, besteht die Möglichkeit, dass Aktiencharts Trendeigenschaften aufweisen. Dies wurde wissenschaftlich zwar noch nicht vollständig beschrieben, allerdings ist die Existenz von Trends vor allem für Markttechniker ein Joker und eine Existenzberechtigung. Murphy (2007) verteidigt die „Technische Analyse“ gegenüber der MEH mit einem durchaus interessanten Gedanken. Nach Murphy kann die Zufälligkeit nämlich nur im negativen Sinn beschrieben werden. Das bedeutet, dass man bisher einfach unfähig war systematische Strukturen in Preisbewegungen wissenschaftlich festzuhalten, was aber nicht bedeutet, dass sie nicht existieren. Auch das Argument, dass vergangene Preisbewegungen womöglich nichts über die Zukunft aussagen können widerstrebt Murphy. Andere Wissenschaften, von der Meteorologie bis hin zur Fundamentalanalyse, bauen fast ausschließlich auf vergangenen Werten auf. Daher kann die „Technische Analyse“ auch etwa unter die beschreibende Statistik fallen, welche überall anerkannt ist.¹⁷

Grundsätzlich folgt die „Technische Analyse“ einer von zwei Denkschulen, wobei die zweite Denkschule die mathematisch-statistische Analyse vertritt. Die mathematisch-statistische Schule wurde oben bereits erläutert. Sie sagt im Wesentlichen, dass Preisbewegungen reiner Zufall sind und einer Normalverteilung folgen. Allerdings werden in diesem Gebiet auch die Entdeckungen von Benoit Mandelbrot behandelt. Dieser stellte bei der Untersuchung von Arbeitseinkommen

¹⁴ Vgl. Thoma (1997), S. 54

¹⁵ Vgl. Thoma (1997), S. 34

¹⁶ Vgl. Hofmann (1973), S. 187

¹⁷ Vgl. Murphy (2007), S. 36f.

fest, dass diese nicht einer Normalverteilung folgten, was wohl bei reinen Zufallsprozessen zu erwarten gewesen wäre. Stattdessen fand er an den Außenflanken der vermeintlichen Normalverteilung breitere Ausläufer bzw. Eintrittswahrscheinlichkeiten. Interessant ist dabei auch, dass diese Art der Verteilung skaleninvariant ist und daher nicht vom Zeitintervall abhängt. Aufbauend auf den Erkenntnissen von Harold Hurst, die oben bereits beschrieben wurden, entwickelte Mandelbrot dann die Theorie der „unendlich andauernden positiven Autokorrelation“, welche auf ein sogenanntes „Langzeitgedächtnis“ bzw. „long memory“ schließen lässt. Später erweiterte er die Brownsche Bewegung um die Möglichkeit einer solchen Autokorrelation und erstellte somit das Konzept der fraktalen Brownschen Bewegung welche in der modernen Wissenschaft heute erforscht wird und in das Gebiet der Chaostheorie fällt.¹⁸

Im Gegensatz zur mathematisch-statistischen Analyse ist die „Technische Analyse“ wenig theoretisch oder mathematisch, sondern viel mehr praxisbezogen. Sie versucht in erster Linie Aktienkursentwicklungen vorherzusagen indem wiederkehrende (geometrische) Muster erkannt werden sollen. Da die Erfahrung über solche Muster nur über die detaillierte Analyse der vergangenen Daten möglich ist, sieht man hier sehr schnell, dass unter Markttechnikern die Vergangenheit einen großen Einfluss auf die Zukunft hat. Unter anderem spiegeln dabei geometrische Muster massenpsychologische Grundeinstellungen der Marktteilnehmer wider, die immer wieder in bestimmten Situationen auftreten können. Das bedeutet also, dass ähnliches Marktgeschehen (in der Zukunft) zu ähnlichem Verhalten der Marktteilnehmer führt.¹⁹

Neben geometrischen Mustern sind vor allem Trendlinien bzw. Trendkanäle im Zuge der „Technischen Analyse“ interessant, weil damit die Autokorrelation in der Kursdynamik erfasst werden kann. Ein Ausbrechen des Kurses kann dann als Indikator zum Kauf oder Verkauf eines Wertpapiers interpretiert werden. Oft wird nicht nur der Chart als Linienverlauf analysiert. Bei der sogenannten „Kerzenchartanalyse“ werden Preisschwankungen eines Tages als Balken dargestellt, welche Muster entstehen lassen, die sich ihrerseits öfter wiederholen.²⁰

Die Grundideen, auf denen die „Technische Analyse“ heute aufbaut, stammen vor allem von Charles Dow, weswegen diese insgesamt als „Dow-Theorie“ bekannt sind. Die wichtigsten sollen im Folgenden kurz skizziert werden. Zunächst war Dow der Meinung, dass die Summe aller Börsentransaktionen das gesamte Wissen der jeweiligen Handelsplätze widerspiegelt. Es müssen an den Indizes also keine Anpassungen vorgenommen werden, da die Märkte alle bekannten Faktoren reflektieren, welche Angebot und Nachfrage beeinflussen. Der zweite Punkt

¹⁸ Vgl. Hafner (2004), S. 43ff.

¹⁹ Vgl. Hafner (2004), S. 40f.

²⁰ Vgl. Hafner (2004), S. 41

der Dow-Theorie besagt, dass es drei Trends in Bezug auf den Markt gibt. Dow definiert etwa einen Aufwärtstrend als ein Muster von steigenden Gipfeln und Tälern und vice versa für Abwärtstrends. In dieser Arbeit wird, wie unten zu sehen ist, ein Trend nach den Grundüberlegungen der Elliott-Wave-Theorie erstellt, welche eine Variante der „Technischen Analyse“ ist. Dabei wird grundsätzlich der Startwert mit dem dritten darauffolgenden Wert verbunden, was eine Trendgerade bzw. Trenderwartung für die Zukunft ergibt. Jedenfalls beschrieb Dow selber die sogenannten „primären“, „sekundären“ und „unbedeutenden“ Trends. Die primären Trends werden dabei etwa mit Ebbe und Flut an den Stränden der Ozeane verglichen und dauern ein oder mehrere Jahre an. Die Korrekturen der primären Trends sind dann die sekundären bzw. mittelfristigen Trends. Dow definiert den Zeitraum dieser Trends allgemein schon genauer, nämlich mit 3 bis 12 Wochen. Dabei gehen die Kurse meistens zwischen 1/3 und 2/3 der vorigen Trendbewegung zurück. Untergeordnete Trends dauern hingegen in der Regel weniger als drei Wochen an und stellen Fluktuationen innerhalb der mittelfristigen Trends dar.²¹

Für Dow sind mit dem dritten Punkt seiner Theorie vor allem die primären Trends interessant, die er sehr praxisbezogen beschreibt. In der Akkumulationsphase kaufen die sehr scharfsinnigen Investoren. Sie können etwaige Hoch- und Tiefpunkte besser abschätzen und erkennen eher, wann ein Markt gewisse Informationen verarbeitet hat. In der öffentlichen Phase steigen dann Trendnachfolger ein, wenn die Wirtschaftsnachrichten etwa von steigenden Kursen berichten. Die scharfsinnigen Investoren verkaufen dann in der Distributionsphase wieder, bevor ihnen andere Marktteilnehmer zuvor kommen. Ein wichtiger Ansatz von Charles Dow ist natürlich auch wann sich oben beschriebene Trends umkehren. Er geht davon aus, dass es wie in der Physik externe Kräfte gibt, welche einen Richtungswechsel verursachen. Heute gibt es eine Vielzahl an Möglichkeiten solche Umkehrsignale auszuforschen. Unter anderem werden dabei Unterstützungslinien, Widerstandslinien, Kursformationen, gleitende Durchschnitte, etc. untersucht. Die einzige Schwierigkeit besteht darin zwischen einer normalen sekundären Korrektur und dem Beginn eines neuen Trends zu unterscheiden.²²

4.1 Kritikpunkte der „Technischen Analyse“

Da nun ein Überblick über die Grundideen der „Technischen Analyse“ der Finanzmärkte gegeben wurde, muss diese mit einem oft verwendeten Argument konfrontiert werden. Es stellt

²¹ Vgl. Murphy (2007), S. 42f.

²² Vgl. Murphy (2007), S. 45f.

sich nämlich allgemein ein wenig die Frage der „Self Fullfilling Prophecy“. Möglicherweise sind nur deshalb Prognoseerfolge zu erkennen, weil bereits eine genügend große Zahl an Investoren an die Charttechniken glaubt. Zusätzlich gehen sie eventuell davon aus, dass nicht nur sie selbst, sondern auch andere Marktteilnehmer ihre Handlungen auf die „Technische Analyse“ stützen. Dennoch legitimiert die große Anzahl an Charttechnikern diese Art der Marktanalyse, zumal nicht nur Kleinanleger diesen Theorien folgen, sondern auch große institutionelle Anleger.²³ Dennoch gibt es in der Wissenschaft einiges an Kritik gegenüber der „Technischen Analyse“. Führende empirische Arbeiten im deutschsprachigen Raum sind dabei vor allem Hofmann (1973) und Herlitz (1975). Hofmann erkannte etwa, dass Investoren mit der „Technischen Analyse“ schlechthin keine signifikant besseren Renditen erzielten als mit einer vergleichbaren „Buy-and-Hold-Strategie“. Die Buy-and-Hold-Strategie sagt dabei eben genau, dass gewisse Assets gemäß der Portfolio-Theorie zwar gekauft werden können, aber dass man durch strategisches Umschichten des Portfolios im Sinne der „Technischen Analyse“ keine Überrenditen erzielt. Ein Problem ist für Hofmann (1973) am Ende das mehr oder weniger subjektive Einzeichnen von Trends und Formationen in den Aktienchart, obwohl die Gültigkeit der Market-Efficiency-Hypothese für ihn zumindest stark eingeschränkt ist. Interessant sind jedoch die Ergebnisse mit gewissen Strategien der „Technischen Analyse“, vor allem jene der gleitenden Durchschnitte bzw. der Filtertechnik. Zum Verständnis soll hier letztere kurz skizziert werden. Betrachtet wird zunächst das Minimum der Wochenschlusskurse in einem Betrachtungszeitraum. Steigt der Kurs einer Aktie über einen bestimmten Filtersatz, der bei Hofmann (1973) zwischen 2,5% und 50% liegt, über das Minimum, dann soll zum nächsten Wochenschlusskurs gekauft werden. Verkauft wird analog, wenn der Kurs der Aktie um denselben Filtersatz unter das Maximum der im Betrachtungszeitraum befindlichen Wochenschlusskurse fällt.²⁴ Bemerkenswert ist also, dass diese Strategien ohne Berücksichtigung von Spesen und auf Bruttobasis vielversprechend zu sein scheinen. Die Ergebnisse auf Nettobasis und mit steigendem Spesensatz relativieren dies jedoch.²⁵ Zu einem ähnlichen Schluss kam auch Herlitz (1975), dessen Analyse auf gleitende Durchschnitte beruhende Indikatoren zunächst vielversprechend aussah, sich allerdings mit Einbeziehung auf Transaktionskosten relativierte. Durch den Einfluss der Vertreter der Market-Efficiency-Hypothese gab es danach eine Zeit lang keine wesentlichen Vorstöße auf dem Gebiet. Erst in den frühen und sehr späten 90er Jahren belegten dann viele Arbeiten, dass die „Technische Analyse“ doch eine gewisse Vorhersagekraft gegenüber der konventionellen „Buy-

²³ Vgl. Hafner (2004), S. 41f.

²⁴ Vgl. Hofmann (1973), S. 88

²⁵ Vgl. Hofmann (1973), S. 188f.

and-Hold-Strategie“ hätte.²⁶ Allerdings haben die Untersuchungen eine besondere Schwierigkeit. Es existiert nämlich eine Vielzahl von Parametern die vor allem ex post spezifiziert werden müssen, was das Argument erlauben würde, sie wären willkürlich gewählt. Wenig konkret sind auch viele Lehrbücher und selbst Murphy (2007) gesteht, dass das Lesen von Charts eine Kunst ist oder zumindest ein hohes Maß an Geschicklichkeit erfordert. So können sich Charttechniker etwa über eine Marktrichtung generell einig sein, aber sie könnten zu unterschiedlichen Zeitpunkten und in anderer Art und Weise Handlungen setzen.²⁷

Es bleiben also Interpretationsspielräume und sollten Strategien der „Technischen Analyse“ nicht erfolgreich sein, müsste das nicht zwangsläufig an der „Technischen Analyse“ selber liegen, sondern einfach an der Unfähigkeit eines Marktteilnehmers Kurse zu lesen und zu interpretieren. Dasselbe gilt für Untersuchungen in der Literatur die meinen, dass man mit dieser Art der Analyse aus Vergangenheitsdaten wohl keine Zukunftsprognosen abgeben kann.

Diesen Mangel an Objektivität wollen Dorfleitner/Klein (2002) bei ihren Untersuchungen zur Prognosegüte der „Technischen Analyse“ möglichst vermeiden. Sie setzen auf Prognosen von anerkannten Fachleuten, die aus vielen Indikatoren abgeleitet werden. Da deren Prognosen zu Analyse Zwecken über längere Zeit regelmäßig verfügbar sein müssen, greifen Dorfleitner/Klein (2002) ein sehr bekanntes deutsches Anlegermagazin auf, das sie namentlich nicht nennen, welches aber über eine große Reichweite verfügt. Die Prognosen der Experten, die sich hauptsächlich auf den Dow-Jones-Industrial-Average-Index und auf den Nikkei-225-Index beziehen, sind dabei verbal und bewegen sich in fünf Kategorien. Die Kategorien der Prognosen umfassen Bewegungen die „abwärts“, „abwärts/seitwärts“, „seitwärts“, „aufwärts/seitwärts“ bzw. „aufwärts“ verlaufen. Im Zuge der Untersuchung werden also die Prognosen aufgenommen und mit den Preisbewegungen nach drei, fünf und acht Wochen verglichen. Dies würde eine sogenannte „Trefferquote“ ergeben, die aber zunächst für sich alleine nicht viel aussagt. Um aussagekräftige Ergebnisse zu bekommen, muss die Trefferquote mit einem Benchmark verglichen werden. Diese Vergleichszahl ist im Wesentlichen beispielsweise die relative Häufigkeit eines klaren Abwärtstrends im Beobachtungszeitraum. Sie sagt also wie oft ein solcher Trend im Vergleich zu anderen Trends tatsächlich vorgekommen ist. Würde man naiv jede Woche auf einen Abwärtstrend setzen, bekäme man genau den Wert des Benchmarks. Hätte die „Technische Analyse“ also eine signifikante Prognosequalität, müsste die Trefferquote den Benchmark jedenfalls übertreffen.²⁸

²⁶ Vgl. Dorfleitner/Klein (2002), S. 1

²⁷ Vgl. Murphy (2007), S. 34f.

²⁸ Vgl. Dorfleitner/Klein (2002), S. 4ff.

Die Untersuchung zeigt aber letztlich ernüchternde Resultate, mit denen eindeutig bewiesen werden kann, dass die Prognosen mit Hilfe der „Technischen Analyse“ keine signifikant besseren Ergebnisse liefern als mit der naiven Prognose. Dorfleitner/Klein (2002) räumen aber ein, dass in ihrer Arbeit bloß auf Indizes und nicht auf Einzelaktien Bezug genommen wurde und dass nicht auszuschließen sei, dass die Prognosequalität der untersuchten Zeitschrift theoretisch mangelhaft sein könnte.²⁹ Somit gibt es zwar starke Hinweise die gegen die „Technische Analyse“ sprechen, allerdings kann man sie wohl nicht ohne Weiteres vollständig als Methode verwerfen, allein schon weil sie in der Praxis großen Anklang findet.

5 Die Trendstrukturanalyse der simulierten Daten

Einer der wichtigsten Streitpunkte zwischen Anhängern der „Technischen Analyse“ und Anhängern der Market-Efficiency-Hypothese ist das Vorhandensein von Trends. In den Ausführungen oben wurden eher längerfristige Trends beschrieben, wobei sich in der Praxis in einem Betrachtungszeitraum der Endwert einer Preisbewegung tendenziell vom Startwert wegbewegt anstatt um diesen herumzuwandern, was, wie bereits erwähnt, eher gegen die MEH spricht. In diesem Teil der vorliegenden Arbeit soll aber vor allem die Struktur von Trends untersucht werden. Wie solche Trends definiert werden wird unten ausführlich beschrieben. Es gilt herauszufinden, wie viele verschieden lange Trendverläufe es in einem bestimmten Betrachtungszeitraum von einem Handelstag gibt und wie oft diese im Durchschnitt an solch einem Tag vorkommen können. Danach stellt sich die Frage ob eine gewisse Gesetzmäßigkeit in der Struktur der Trends festgestellt werden kann. Würde die Struktur von Trendverläufen beispielsweise einer theoretischen Verteilung folgen, könnte man in weiterer Folge etwa eine Formel entwickeln um kürzer- oder längerfristige Trendverläufe statistisch vorherzusagen. Die Entwicklung einer solchen Gesetzmäßigkeit würde zwar den Rahmen dieser Arbeit sprengen, dennoch kann die statistische Analyse der Struktur von Trends eine Grundlage dazu liefern. Da in der Praxis vor allem die geometrische Brownsche Bewegung benutzt wird um Aktienverläufe zu simulieren, wird diese zunächst einer Analyse unterworfen. Zur Programmierung wird dabei das freie Programm „GNU R“ verwendet. Jedenfalls wird anschließend derselbe Algorithmus auf echte historische Kursdaten des EURO STOXX 50 angewendet. Diese werden natürlich auch einer Analyse der Trendstruktur unterworfen. Besonders interessant ist allerdings im Anschluss der Vergleich der simulierten bzw. der realen Werte. Die Frage die sich dabei stellt ist, ob sich

²⁹ Vgl. Dorfleitner/Klein (2002), S. 21f.

die Struktur der Anzahl der verschiedenen langen Trends der simulierten Daten von jenen der realen Daten unterscheidet. Da die Brownsche Bewegung als Weiterentwicklung des Random Walk zumindest für Charttechniker die Realität nicht gut genug beschreibt, könnte man annehmen, dass die Struktur von Trends beider Varianten also verschieden ist. Wäre sie jedoch sehr ähnlich, könnte man möglicherweise annehmen, dass die Brownsche Bewegung die Realität doch gut beschreibt, was wohl bedeuten würde, dass mehr Zufall in Preisbewegungen steckt, als angenommen.

Ausgangspunkt zur Analyse von etwaigen Trends in Indizes stellt zunächst die Simulation solcher Preisbewegungen dar. In der Praxis können durch genau solche Simulationen unter anderem Aktienwerte einigermaßen gut geplant werden um wiederum etwa den Wert von Derivaten zu bestimmen.³⁰

Für diesen Zweck ist die Brownsche Bewegung besonders geeignet, da sie der bekannteste zeitkontinuierliche stochastische Prozess ist und zuerst historisch untersucht bzw. später mathematisch modelliert wurde. Diese beruht auf den Entdeckungen eines schottischen Botanikers namens Robert Brown im Jahr 1828. Brown fand Zusammenstöße kleinster Teilchen auf Flüssigkeitstropfen, die eine sehr unregelmäßige Bewegung erzeugten. Die mathematische Ausformulierung fand aber erst später durch Norbert Wiener statt, weswegen die Brownsche Bewegung auch oft als „Wiener Prozess“ bezeichnet wird.³¹ Während Hafner (2004) allerdings die Brownsche Bewegung, wie sie heute benutzt wird, als Zufallsprozess beschreibt, wird dies von Thoma (1997) relativiert. Ihm zufolge nimmt man zwar an, dass Browns Blütenteilchen zufällig aufeinanderstoßen und somit einer Normalverteilung folgen, aber hundertprozentig beweisen könnte man diese Annahme bisher nicht.³²

Die einzelnen Zuwächse der Brownschen Bewegung sind also identisch verteilt und nicht miteinander korreliert.³³ Die Zuwächse sehen formal wie folgt aus:

$$W_{t_{i+1}} = W_{t_i} + \sqrt{t_{i+1} - t_i} * Z_i$$

Die Zufallsvariable Z_i bzw. die Änderungen der Brownschen Bewegung sind normalverteilt, weswegen sie aber auch negative Werte annehmen können. Da dieser Umstand offensichtlich, gerade für Aktienindizes, unrealistisch ist, wird also konkret die geometrische Brownsche

³⁰ Vgl. Griesing (2012), S.3

³¹ Vgl. Hafner (2004), S. 43

³² Vgl. Thoma (1997), S. 17

³³ Vgl. Hafner (2004), S. 47

Bewegung verwendet, welche eine Modifikation der Brownschen Bewegung W_t ist und sowohl multiplikative als auch unabhängige Zuwächse hat.³⁴

Es wird angenommen, dass Aktienkurse einer geometrischen Brownschen Bewegung folgen:

$$S(t) = x * \exp \left[\left(r - \frac{\sigma^2}{2} \right) t + \sigma W(t) \right]$$

Zu beachten ist, dass vor allem der Startwert x , die Volatilität σ und der Drift r Konstanten darstellen die angenommen bzw. festgelegt werden. Die Volatilität gibt als Risikomaß an wie stark der Ausschlag der Brownschen Bewegung ist und wie sehr sie sich an die x-Achse anschmiegt. Nach Griesing (2012) nähert sich diese der x-Achse stärker an und hat größere Schwankungen, wenn die Volatilität höher ist.³⁵ In dieser Arbeit nimmt die Volatilität der simulierten Brownschen Bewegung den Wert 1 an, wie in Schritt 1 im Abschnitt 1 des Anhangs zu erkennen ist.

Wie bereits erwähnt, beeinflusst auch der Drift die Ausprägung der simulierten Kursverläufe. Dieser Parameter bestimmt, ob sich der Pfad der Bewegung tendenziell nach oben bzw. nach unten bewegt was entweder durch einen positiven oder negativen Drift initiiert wird.³⁶ Da in dieser Arbeit Trendverläufe analysiert werden sollen, wird versucht im Vorhinein keine Bewegungstendenz in die Brownsche Bewegung einfließen zulassen, weswegen der Drift 0 gesetzt wird.

Eine andere Konstante stellt die Anzahl der Punkte bzw. Preise N in dem zu analysierenden Index dar. Konkret wurden in dieser Arbeit 840 Punkte pro simulierte Bewegung gewählt. Der Grund dafür ist, dass im weiteren Verlauf der hier verwendete Algorithmus auf echte historische Kurse des EURO STOXX 50 angewendet wird. Um später also einen Vergleich zu ermöglichen, wird hier an der Anzahl der Punkte des EURO STOXX festgehalten. Jeder dieser Handelstage beinhaltet 50400 Punkte bzw. Ticks. Ein Tick ist dabei die kleinst mögliche Kursänderung in einem Aktienchart. Da mit einer Vielzahl von verschiedenen Handelstagen gearbeitet wird, wäre die Verwendung von 50400 Punkten pro Tag technisch unnötig schwierig. Nimmt man an, dass etwa ein Tick pro Sekunde entsteht, reicht es vollkommen aus dem Datenvektor eines Handelstages von Sekunden auf Minuten zu kürzen, was also die 840 Punkte der simulierten Preisbewegungen erklärt. Wie weiter in Schritt 1 zu sehen ist, wird technisch zuerst ein leerer

³⁴ Vgl. Griesing (2012), S.4

³⁵ Vgl. Griesing (2012), S.7

³⁶ Vgl. Griesing (2012), S.7f.

Vektor W erstellt, welcher dann zunächst durch Ausprägungen der Brownschen Bewegung ersetzt wird. Danach werden diese Ergebnisse, wie oben beschrieben, durch Werte der geometrischen Brownschen Bewegung S adaptiert.

Die Grundüberlegungen zur Analyse der (simulierten) Aktienindizes im Anschluss stützen sich teilweise, wie bereits oben erwähnt, auf die Theorien von Ralph Nelson Elliott, welche eine Variante der klassischen „Technischen Analyse“ darstellen. Dieser entwickelte in den späten 1920er Jahren die Elliott Wave Analyse zur „Technischen Analyse“ bzw. zur etwaigen Vorhersage von Finanzmärkten. Im Kern steht dabei ein natürlicher Rhythmus welcher vor allem Grundverhaltensmuster in Bezug auf die Psyche bzw. Emotionen wie Angst, Gier, Optimismus oder Pessimismus hinsichtlich der Entwicklung von Märkten widerspiegelt.³⁷ Bei der Elliot Wave Analyse, als auch bei den Analysen in vorliegender Arbeit, werden ausschließlich das Kursverhalten eines Marktes in der Vergangenheit bewertet.³⁸

Dabei geht man von zwei Grundmustern aus mit denen sich Marktverhalten beschreiben lässt. Zum einen spricht man also von Antriebs- bzw. Impulsbewegungen, die praktisch Aufwärtsbewegungen darstellen. Zum anderen spricht man von Korrektur- bzw. Abwärtsbewegungen. Eine „Welle“ besteht nach Elliott aus fünf Impulswellen (-bewegungen) und drei Korrekturwellen (-bewegungen), welche aber auch auf einzelne Auf- bzw. Abbewegungen heruntergebrochen werden kann. So kann bereits ein einzelner Preissprung nach oben eine Impulswelle bzw. ein Preissprung nach unten eine Korrekturwelle darstellen.³⁹ Um anschließend eine Trendlinie ziehen zu können, müssen zunächst Anzeichen vorhanden sein. Nach Murphy (2007) ist dies beim positiven Trend frühestens der Fall, wenn es mindestens zwei „Reaktionstiefs“ gibt, von denen eines höher liegt als das andere.⁴⁰

5.1 Der Algorithmus zur Untersuchung von Trendstrukturen in simulierten Daten

Diese Arbeit folgt im Wesentlichen der Idee von Impuls- und Korrekturwellen nach Elliott bzw. Murphy, wonach also zwei einzelne Bewegungen einen Trendverlauf für die Zukunft indizieren. Dabei entstehen acht (2^3) Möglichkeiten die entweder in einem Trend mit positivem Anstieg oder mit negativem Anstieg resultieren.

³⁷ Vgl. Heussinger (2002), S. 1

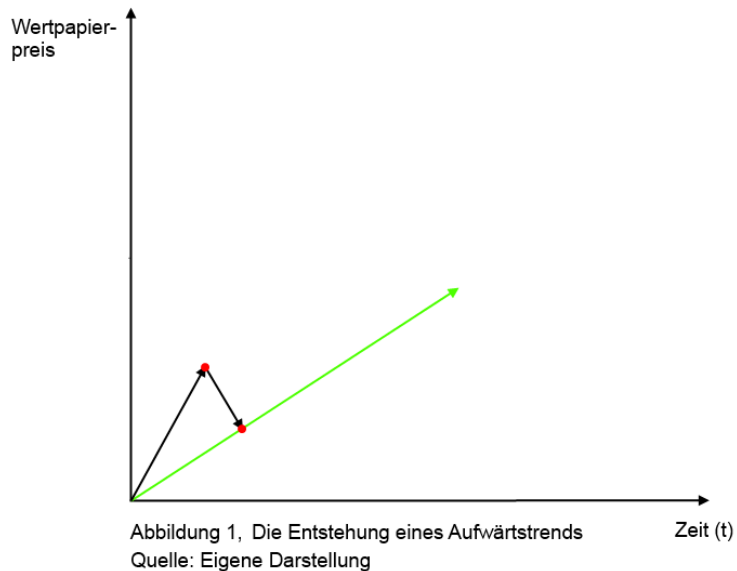
³⁸ Vgl. Heussinger (2002), S. 3

³⁹ Vgl. Heussinger (2002), S. 24

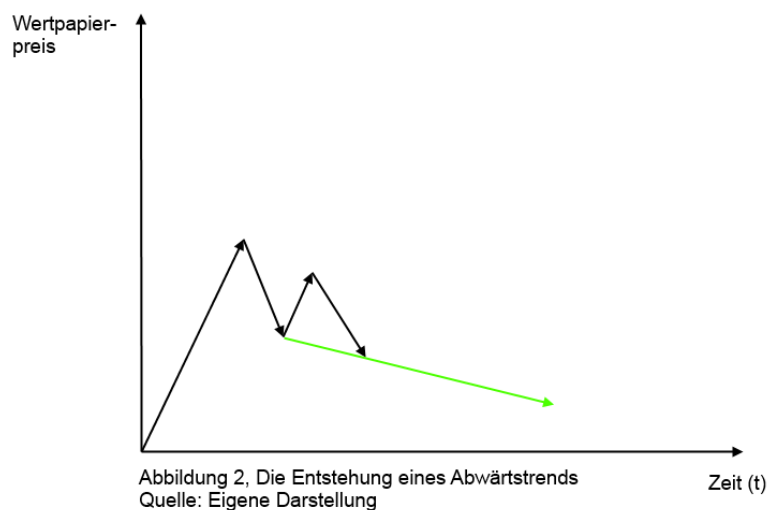
⁴⁰ Vgl. Murphy (2007), S.78

5.1.1 Die Erzeugung von Trendgeraden

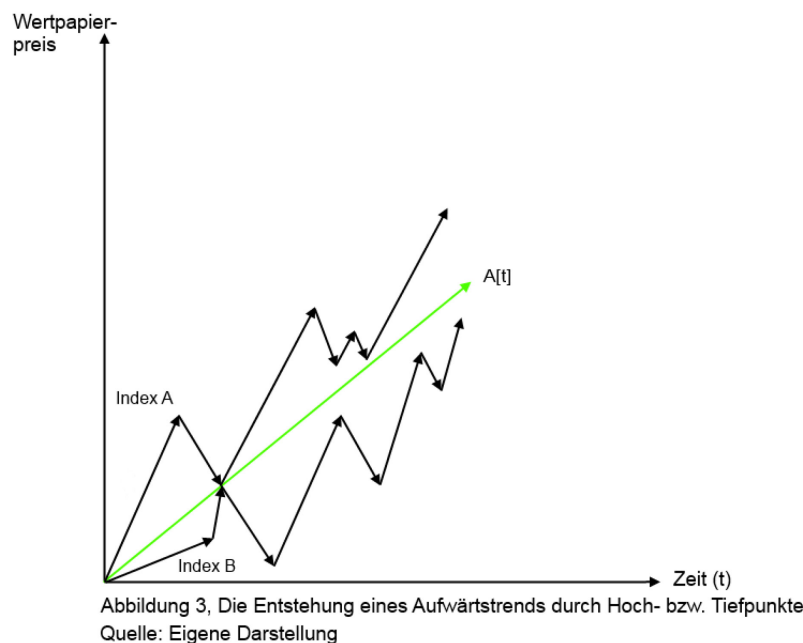
Folgt auf eine bestimmte positive Preisbewegung eine kleinere Korrekturbewegung ergibt sich, wie in Abbildung 1 zusehen ist, ein Trend mit positiver Steigung.



Umgekehrt könnte auch eine Abwärtsbewegung und eine darauf folgende größere Aufwärtsbewegung, in derselben Dimension wie in Abbildung 1, einen positiven Trend initiieren. Selbstverständlich resultiert ein steigender Trend auch aus zwei positiven Impulsen, wobei ein flacherer Anstieg auf einen steileren folgt oder umgekehrt.



Nach dieser Logik entsteht ein Trend mit negativer Steigung durch eine Impulsbewegung und eine darauf folgende größere Korrekturbewegung, wie in Abbildung 2, bzw. durch einen größeren negativen Preissprung, gefolgt von einem kleineren positiven Sprung. Wie beim positiven Trend kann natürlich auch ein negativer Trend aus zwei negativen Preissprüngen entstehen, wobei einer dieser Sprünge für gewöhnlich steiler bzw. flacher ist als der andere.



Wie bereits oben erwähnt, wird im Zuge dieser Arbeit unterschieden, ob ein Aufwärtstrend durch Hochpunkte oder durch Tiefpunkte eingeleitet wird. Betrachtet man Abbildung 3, bedeutet das, dass hier konkret unterschieden wird, ob der vierte Punkt ab dem Startpunkt der Trendgeraden über oder unter dieser liegt. Diese Unterscheidung ist vor allem bei der Ablöse „alter“ Trends durch neue wichtig, was weiter unten genauer erläutert wird.

Wie in Abbildung 3 also zu erkennen ist, stellen die Bewegungen A und B beide einen Trend nach oben dar. Index A beginnt mit einer steilen Impulswelle und einer kleineren Korrekturwelle, während Index B mit zwei aufeinander folgenden Aufwärtsbewegungen beginnt. Technisch ergibt sich also aus den relevanten ersten und dritten Punkten innerhalb eines simulierten Index nach der allgemeinen Geradengleichung $A = k * x + d$ eine Trendgerade $A[t]$, wobei $t=0$ für den jeweiligen Startzeitpunkt steht. Programmiertechnisch ist hier eine kleine Besonderheit zu beachten. Die Gerade $A[t]$ besteht aus einzelnen Punkten und zwar zu jedem Zeitpunkt t_i zu dem per Definition auch ein Preis des Aktienindex existiert. Da, wie oben bereits erwähnt, eine Trendlinie aus dem ersten und dritten Punkt eines Kurses besteht, müssen vor allem der jeweils dritte Punkt (gemessen am Startwert der jeweiligen Gerden) des Aktienkurses und der dritte

Punkt der Trendgeraden $A[t]$ selber identisch sein. Allerdings rechnet GNU R mit maximal 16 Nachkommastellen, wobei der Rest aus dem sogenannten „noise“ besteht und an der maximal verwendbaren Stelle gerundet wird. Aus diesem Grund erkennt R diese zwei Zahlen nicht als gleiche Zahlen, obwohl sie das sein müssten. Daher besteht technisch die Notwendigkeit die zwei betreffenden Zahlen manuell als „gleich“ zu definieren, wie das in Abschnitt 1/Schritt 5 im Anhang zu sehen ist.

5.1.2 Die Initiierung neuer Trends bei Durchbrechen von Preisbewegungen

Eines der elementarsten Themen dieser Arbeit ist, wie bzw. wann ein entstandener Trend endet und wie dieser dann von einem neuen Trend abgelöst wird. Dazu beschreibt Murphy (2007) vor allem den Bruch der Trendlinie als Signal zur Trendumkehr.⁴¹

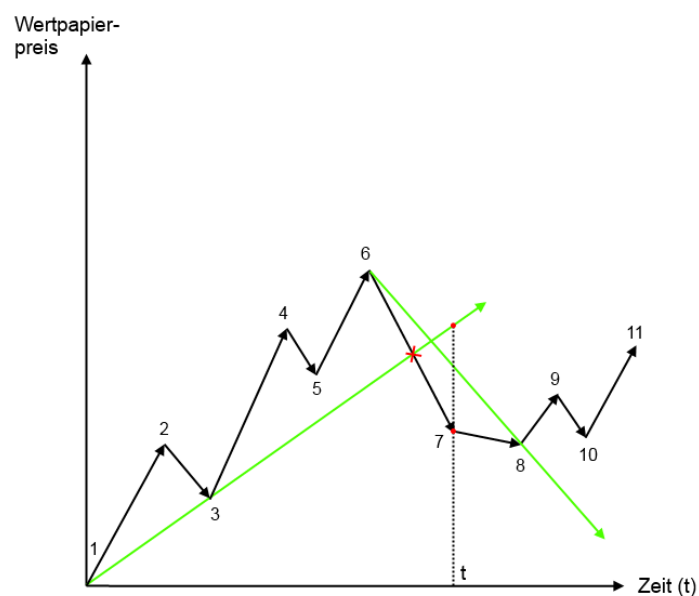


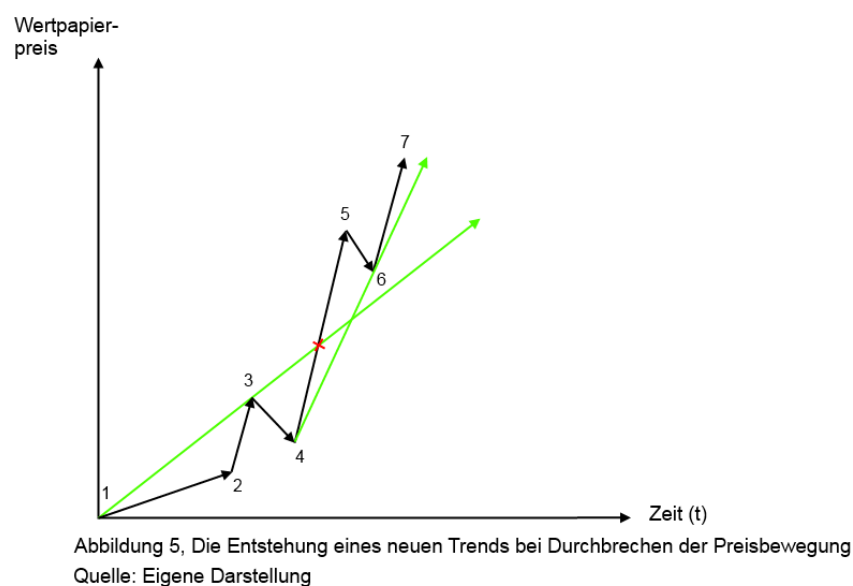
Abbildung 4, Die Entstehung eines neuen Trends bei Durchbrechen der Preisbewegung
Quelle: Eigene Darstellung

Wie in Abbildung 4 zu sehen ist, entsteht also die erste aufsteigende Trendgerade aus Punkt 1 und Punkt 3. Durchbrochen wird diese Gerade von einer stärkeren Abwärtsbewegung zwischen Punkt 6 und Punkt 7. Der „Bruch“ stellt sich praktisch so dar, dass Punkt 7 zu genau diesem Zeitpunkt t tiefer liegt als der zeitgleiche Punkt auf der ersten Trendlinie $A[t]$, d.h. $S[7] <$

⁴¹ Vgl. Murphy (2007), S.81

$A[[1]][7]$. Dies initiiert in diesem Fall einen (kurzfristigen) Abwärtstrend dessen relativer erster Punkt bei Punkt 6 liegt und durch den relativen dritten Punkt (Punkt 8) verläuft. Nach derselben

Logik würde ein entsprechender erster Abwärtstrend durch einen Bruch einen neuen Aufwärtstrend einleiten. Von besonderer Bedeutung ist allerdings, wo sich der vierte Punkt der Brownschen Bewegung ab Anfang der Trendlinie befindet. Befindet sich dieser, wie in Abbildung 5 (anders als in Abbildung 4), unterhalb der aktuellen Trendgeraden $A[t]$, kann die Preisbewegung nur dann von dieser durchbrochen werden, wenn frühestens Punkt 5, zu $t=5$, über der Geraden liegt.

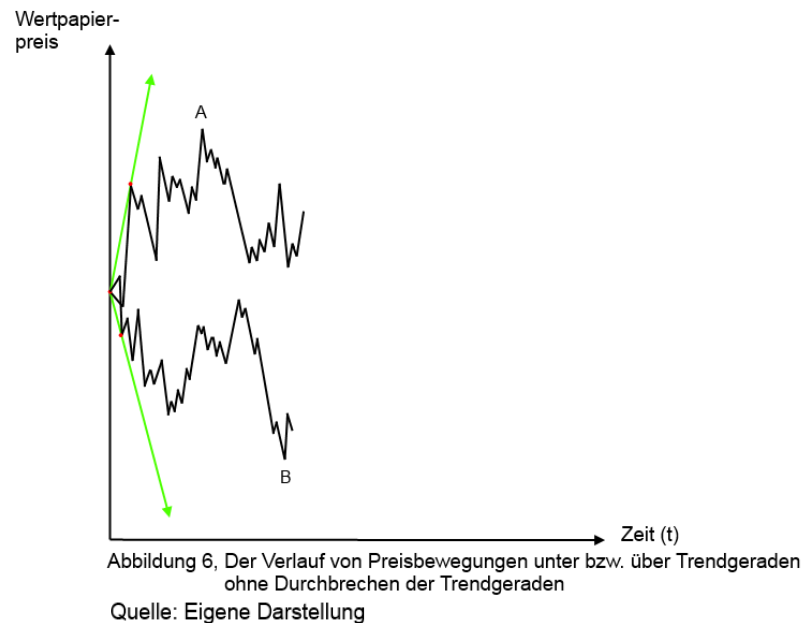


Weiters ist zu erkennen, dass ein Bruch eines Aufwärtstrends mit einem Punkt 4, der unter diesem Trend liegt, einen neuen Trend erzeugt, welcher oftmals auch aufsteigend ist. Dasselbe gilt für fallende Trendgeraden, deren korrespondierender vierter Index-Punkt über der betreffenden Geraden liegt.

5.1.3 Die Initiierung neuer Trends ohne Durchbrechen von Preisbewegungen

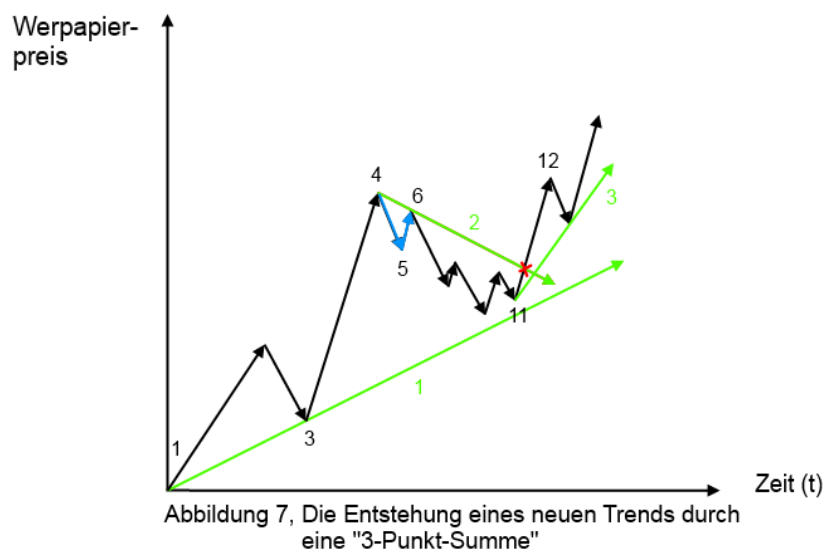
In einigen Fällen kommt es gar nicht erst zu einer Ablöse eines Trends durch einen anderen, da die Preisbewegung in weiterer Folge nicht mehr den Verlauf der Trendlinie kreuzt.

Besonders deutlich fällt dies auf, wenn die Steigung k sehr hoch bzw. sehr niedrig ist, d.h. wenn die Trendgerade sehr steil verläuft.



In Abbildung 6 ist gut zu erkennen, dass beide Start-Trends von A und B zunächst von relativ großen Kurssprüngen eingeleitet werden und die jeweiligen Indexbewegungen anschließend auf der Unter- bzw. Oberseite dieser Geraden weiterverlaufen. Offensichtlich gibt es aber bei beiden Preisentwicklungen kurz- und mittelfristige Auf- bzw. Abwärtsbewegungen, welche an diesem Punkt unbeachtet bleiben. Mit der Zeit wird es sogar so gut wie unmöglich, dass selbst längere und steilere Bewegungen in Richtung der Trendlinien diese auch tatsächlich durchbrechen können.

Da allerdings auch in solchen Fällen offensichtliche Trends erfasst werden sollen, wird ein weiteres Abbruchkriterium benötigt. Daher wird auf die grundsätzliche Idee der Trendentstehung zurückgegriffen, wonach etwa ein fallender Trendverlauf zumindest durch einen kurzfristigen Trendanstieg abgelöst wird oder umgekehrt.



Im Wesentlichen wird daher, wie dies in Abbildung 7 zu sehen ist, ab Punkt 3 nach Indizien gesucht welche, in dem Fall, einen fallenden Trend einleiten könnten. Konkret wird also, wie oben bereits beschrieben, nach einer Korrekturwelle gesucht, deren nachfolgende Bewegung eine kleinere Impulswelle darstellt. Für die technische Vorgehensweise ist diesbezüglich Abschnitt 1/Schritt 3 im Anhang relevant. Der Algorithmus berechnet im Vorhinein die Summen aus allen drei aufeinanderfolgenden Punkten, d.h. dass etwa der Betrag des Anstiegs zwischen Punkt 1 und Punkt 2 zu dem Betrag des Abstiegs zwischen Punkt 2 und Punkt 3 addiert wird. Dasselbe gilt für Punkt 2,3 und 4, Punkt 3,4 und 5 etc. Diese Summen sind in den allermeisten Fällen positiv oder negativ und werden in einem Vektor abgespeichert. Im konkreten Fall in Abbildung 6 leitet Punkt 4 einen neuen Abwärtstrend ein, da hier die Impulswelle zum ersten Mal kleiner ist als die Korrekturwelle, d.h. die Summe der beiden Bewegungen ist negativ. Wie bereits oben erläutert wurde, initiieren Punkt 4 und Punkt 6 also einen neuen Abwärtstrend, welcher seinerseits durch einen Trendbruch zwischen Punkt 11 und Punkt 12 von einem neuen Trend abgelöst wird.

5.1.4 Die Erzeugung sämtlicher Trendgeraden innerhalb einer geometrischen Brownschen Bewegung

Aus programmiertechnischer Sicht ist für diese Prozedur Abschnitt 1/Schritt 5 im Anhang zu beachten. Zunächst werden der Iterationsvariable i und einer Dummy-Variable h jeweils der Startwert 1 zugeordnet. Danach findet der Sprung in die While-Schleife statt. Mit dem Startwert von i wird zu Beginn die erste Gerade durch den ersten und dritten Punkt der simulierten Brownschen Bewegung gelegt. Danach wird, wie oben bereits beschrieben, nach dem Abbruchpunkt h gesucht. Dabei muss der Algorithmus unterscheiden ob es sich zunächst bei der ersten Trendlinie um eine Linie mit positiver oder negativer Steigung handelt, ob Punkt 4 der Index-Bewegung über bzw. unter dem zeitgleichen Punkt von $A[t]$ liegt und ob die Trendlinie im weiteren Verlauf den Index schneidet oder nicht. Daraus ergeben sich, wie in Abschnitt 1/Schritt 5 im Anhang zu sehen ist, acht einzelne Möglichkeiten die jeweils zu überprüfen sind. Nur wenn alle drei Argumente innerhalb einer (if-) Bedingung zutreffen, wird der Abbruchpunkt h und in weiterer Folge der neue Startpunkt i ermittelt. Dieser Vorgang wiederholt sich natürlich nur solange bis das Ende der simulierten Preisentwicklung erreicht ist.

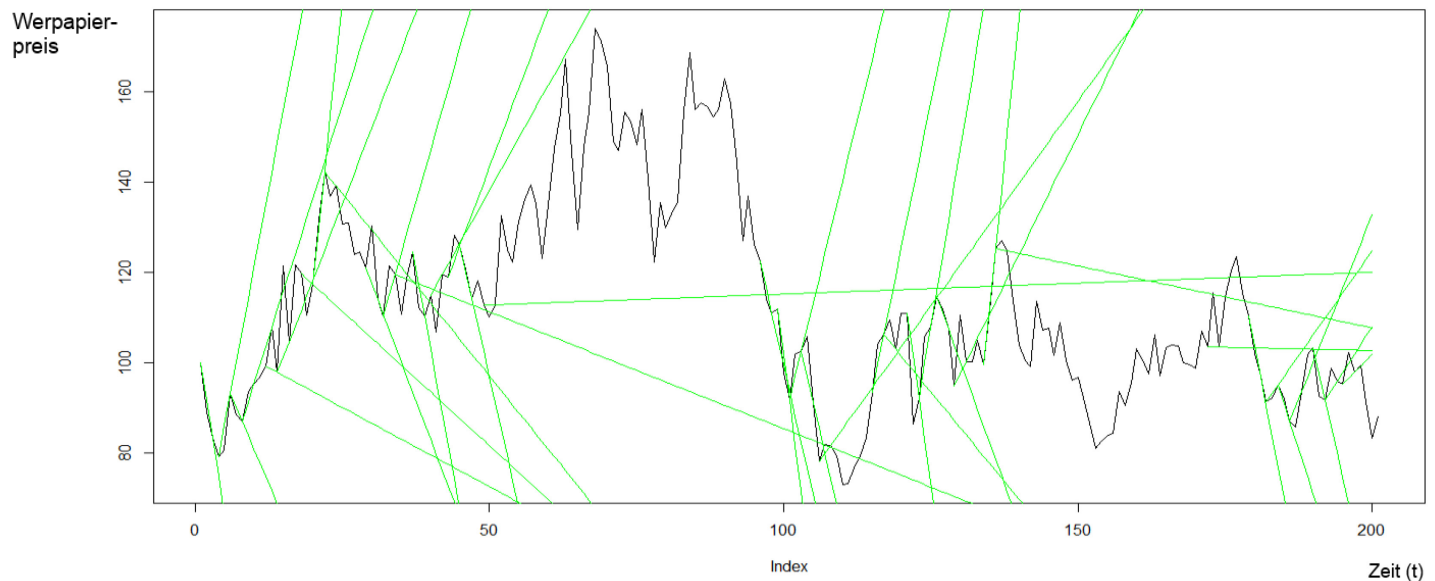


Abbildung 8, Beispiel einer simulierten Preisbewegung über 200 Ticks bzw. Minuten mit sämtlichen Trendverläufen
 Quelle: Eigene Darstellung

In Abbildung 8 ist dazu ein Beispiel einer geometrischen Brownschen Bewegung mit deren Trendverläufen zu sehen. Aus Übersichtsgründen wurden allerdings nur 200 mögliche Punkte im Index vorgesehen.

6 Statistische Auswertung der simulierten Daten

Wie bereits oben beschrieben, werden durch den erzeugenden Algorithmus mit der entsprechenden Vorgangsweise Trendgeraden in eine simulierte, geometrische Brownsche Bewegung mit 840 Punkten gelegt. Der kleinste Trend umfasst per Definition zwei Punkte. Der erste Punkt ist dabei der auf den Startwert der Geraden folgenden Punkt. Der zweite Punkt fällt mit dem dritten Punkt ab dem Startwert und somit mit dem zweiten erzeugenden Punkt der Trendgeraden zusammen. In allen Fällen ist ein etwaiges Abbruchkriterium von Belang. Erst, wenn eine Trendgerade durch eine nächste abgelöst wird, lässt sich die Länge der vorhergehenden Trendgeraden bestimmen. Die Länge dieser Trends definiert sich also über die Anzahl der in diesem Trend befindlichen Punkte vom zweiten Punkt ab dem Startwert bis (inklusive) zum Startwert des neuen Trends. Grundsätzlich könnte man auch den ersten Punkt der Trendgeraden als Teil des Trends annehmen, allerdings geht von nur einem Punkt, ohne vorhergehende Bewegungen, noch kein Trendcharakter aus, weshalb er also nicht mitgezählt

wird. Allerdings ist hier zu bemerken, dass es sich bei dieser Überlegung tatsächlich um reine Definitionssache handelt.

Weshalb der erste Punkt des neuen Trends noch zum vorherigen Trend zählt, ist beispielsweise in Abbildung 5 ersichtlich. Der erste Trend definiert sich über eine flachere Bewegung unterhalb der Trendgeraden. Da der erste Punkt des neuen Trends noch dazu passt, wird er also dem alten Trend zugerechnet.

Somit entstehen in einer zunächst simulierten Bewegung Trends mit unterschiedlicher Länge, wobei auch Trends mit gleicher Länge öfter vorkommen. Im weiteren Verlauf ist es daher auch von Interesse, wie viele sehr kurzfristige Trends (d.h. Trends über zwei oder drei Punkte) bzw. etwas längere Trends in einem Tageschart zu finden sind.

Da man aber mit bloß einer einzelnen Simulation bzw. Trendanalyse wenig Aussagemöglichkeiten hat, ist es von Vorteil Preisbewegungen öfter zu simulieren um dann etwa zu bestimmen, wie viele 2er-, 3er-, 4er-Trends, etc. im Durchschnitt in einem Chart vorherrschend sind. Die Verteilung, die sich z.B. aus 10 000 Simulationen ergibt, wird im Folgenden „Masterverteilung“ genannt. Die Anzahl der Iterationen ist dabei frei wählbar. Da, wie später noch genauer erläutert wird, die vorliegenden historischen Daten des EURO STOXX genau 2131 Handelstage (das entspricht den Tagen von 2005 bis 2013) umfassen und ein etwaiger Vergleich möglich sein soll, wird eine simulierte „Masterverteilung“ in dieser Arbeit aus 2131 Iterationen erzeugt. Die Vorgangsweise wird im Anhang beschrieben.

Zunächst wird bei jeder Iteration des Algorithmus ein Vektor mit dem Namen „liste“ erzeugt. Dieser Vektor enthält alle Startpunkte von Trendgeraden innerhalb einer simulierten, geometrischen Brownschen Bewegung. Da der erste Trend per Definition immer mit dem ersten Preis/Punkt im Chart anfängt, beträgt der erste Eintrag in „liste“ immer „1“. Mit jeder Wiederholung des Algorithmus wird auch immer ein neuer rein zufälliger Bewegungsablauf erzeugt. Damit entstehen auch immer unterschiedlich lange Trends welche auch unterschiedlich oft im jeweiligen Chart vorkommen. Im Wesentlichen werden alle Vektoren „liste“ in einer R-Liste abgespeichert und anschließend in „Differenz-Vektoren“ umgewandelt. Dabei werden jeweils Differenzen zwischen einem Element und dem vorhergehenden Element gebildet. Somit enthalten die neuen Vektoren, genannt „trendtestdiff“, die jeweiligen Längen der Trends im dazugehörigen Tageschart. Die Trends werden anschließend zusammengefasst, nach ihrer Größe geordnet und es wird ein Durchschnitt über 2131 Iterationen berechnet.

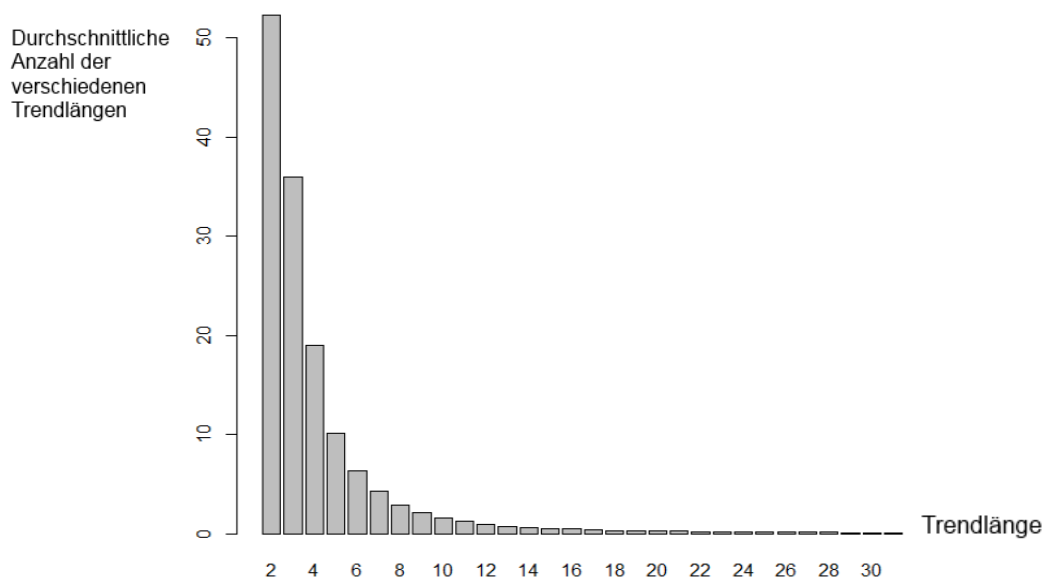


Abbildung 9, Die Trendstruktur der "Masterverteilung" der simulierten Daten
Quelle: Eigene Darstellung

Abbildung 9 zeigt das Ergebnis der 2131 Simulationen. Offensichtlich dominieren vor allem kurze Trends welche nur zwei, drei oder vier Punkte enthalten. Konkret gibt es unter den Simulationen im Durchschnitt ca. 450 Trends auf 840 Punkte im Chart, wobei aus Übersichtsgründen in Abbildung 9 nur 30 verschiedene Trendlängen zu sehen sind. Dieser Mittelwert ist aber nicht allzu streng zu interpretieren, da sehr viele und vor allem längere Trends nur ein oder zwei Mal unter allen Iterationen vorkommen. Trends die nur zwei Punkte umfassen, kommen im Schnitt rund 52 Mal vor, während Trends die drei Punkte umfassen, durchschnittlich in einem Tageschart nur mehr rund 36 Mal vorkommen. Wie in Abbildung 8 zu sehen ist, fällt die Häufigkeitskurve sehr schnell ab und so bewegen sich die Trendlängen welche im Schnitt mindestens einmal pro simulierter Brownscher Bewegung vorkommen zwischen zwei und elf bzw. zwölf. Aus der Erfahrung mit weniger als rund 2000 Iterationen zeigt sich, dass sehr große Trends erst bei mehreren Iterationen vorkommen. Bei mehr als 2000 Iterationen kommen dann auch bald alle Trendlängen zwischen 2 und 840 Punkten in den Simulationen vor, was statistisch recht leicht nachzuvollziehen ist.

Wie in Abbildung 8 deutlich zu sehen ist, ergibt die stetig fallende Anzahl länger andauernder Trends eine sehr schöne Kurve, wenngleich sie auch sehr rasch gegen die x-Achse läuft. Um herauszufinden ob die durchschnittliche Anzahl von Trends in einem Chart einer theoretischen Verteilung folgt, wird in der Statistik ein Anpassungstest verwendet.⁴² Vorher muss allerdings

⁴² Vgl. Stahel (2008), S. 247

bestimmt werden, welche theoretische Verteilung gegen eine bestimmte Stichprobe getestet werden soll. Im konkreten Fall soll also getestet werden ob die Anzahl der verschiedenen Trendlängen einer geometrischen Verteilung folgt. Diese eignet sich hier, weil sie vor allem eine diskrete Verteilung darstellt und etwa eine Lebensdauerverteilung bzw. Wartezeitenverteilung modelliert.⁴³ So kann in dieser Arbeit die Anzahl der Punkte (d.h. die Trendlänge) innerhalb eines Trendverlaufs als „Lebensdauer“ dieses Trends interpretiert werden.

Als nächster Schritt gilt es zu bestimmen, welcher Anpassungstest überhaupt verwendet werden soll. Hilfreich sind dabei „Entscheidungsbäume“ welche je nach Datenlage über den passenden Test informieren. Da die Anzahl der jeweiligen Trends metrische Daten sind und wir einen Test auf eine diskrete Verteilung durchführen wollen, eignet sich hier der Chi-Quadrat-Anpassungstest.⁴⁴ Dabei lauten die Null- bzw. Alternativhypothese wie folgt:

H_0 : *Die zu testende Verteilung zeigt keine großen Unterschiede zu den Erwartungen einer geometrischen Verteilung.*

H_1 : *Die zu testende Verteilung zeigt signifikante Unterschiede zu den Erwartungen einer geometrischen Verteilung.*

Die technische Vorgangsweise wird im Anhang im Abschnitt 2 erläutert. Grundsätzlich werden die Daten (Anzahl der jeweiligen Trendlängen) in Klassen unterteilt, mit einem Parameter, dem P-Wert, geometrisch verteilt und anschließend mit den tatsächlichen relativen Häufigkeiten verglichen. Dabei ist technisch zu beachten, dass die geometrisch verteilten Häufigkeiten in Summe 100 Prozent (d.h. „Eins“) ergeben. Da allerdings, wie oben bereits erwähnt, in R bei Berechnungen Ungenauigkeiten („Noise“) auftreten, muss der Anteil der letzten Klasse die Differenz aus „Eins“ und der Summe der Anteile aller anderen Klassen sein.⁴⁵

Der P-Wert ist dabei, ähnlich wie „Lambda“, der Kehrwert des Mittelwerts über die Daten. Im konkreten Fall dieser Arbeit wird der Mittelwert aus den multiplizierten Werten zwischen der Trendlänge an sich und den jeweiligen Häufigkeiten gebildet. Allerdings muss man der Vollständigkeit halber noch eine kleine Anpassung vornehmen. Per Definition umfasst der kleinste Trendverlauf zwei Punkte. Bei einer geometrischen Verteilung muss der Startwert aber „Null“ sein. Aus diesem Grund müssen die Werte von „Trendlänge“ noch vor der Multiplikation

⁴³ Vgl. Kohn/Öztürk (2010), S. 177

⁴⁴ Vgl. Universität Marburg (1998), [Zugriff am 13.08.2015]

⁴⁵ Vgl. Hatzinger/Hornik/Nagel (2011), S. 266

mit ihren Häufigkeiten um zwei Einheiten nach links verschoben werden, wie im Anhang zu sehen ist.

Des Weiteren werden für den Chi-Quadrat-Test noch die relativen Häufigkeiten der Daten benötigt, was allerdings recht einfach ist. Dafür wird einfach die jeweilige Anzahl der Trends durch die Gesamtzahl aller Trends dividiert. Analog zu Hatzinger/Hornik/Nagel (2011) werden noch die erwarteten Häufigkeiten berechnet. Diese sind das Produkt aus den Werten der geometrischen Verteilung und der Anzahl aller verschiedenen Trendklassen bzw. – längen.⁴⁶

Für die eigentliche Auswertung mittels Chi-Quadrat-Anpassungstest muss im Wesentlichen eine Sichtung der Daten erfolgen. Einen Auszug der berechneten Werte zeigt die nachfolgende Tabelle 1.

	absolut	relativ	geometrisch	erwartete
[1,]	0	3.608868e-01	2.694393e-01	2.185153e+02
[2,]	1	2.485370e-01	1.968418e-01	1.596387e+02
[3,]	2	1.307028e-01	1.438049e-01	1.166257e+02
[4,]	3	6.937840e-02	1.050582e-01	8.520218e+01
[5,]	4	4.395544e-02	7.675137e-02	6.224536e+01
[6,]	5	2.896433e-02	5.607154e-02	4.547402e+01
[7,]	6	2.011207e-02	4.096366e-02	3.322153e+01
[8,]	7	1.455134e-02	2.992644e-02	2.427034e+01
[9,]	8	1.087805e-02	2.186308e-02	1.773096e+01
[10,]	9	8.415020e-03	1.597231e-02	1.295354e+01
[11,]	10	6.696532e-03	1.166874e-02	9.463348e+00
[12,]	11	5.356255e-03	8.524723e-03	6.913550e+00
[13,]	12	4.459230e-03	6.227828e-03	5.050768e+00
[14,]	13	3.771808e-03	4.549806e-03	3.689893e+00
[15,]	14	3.228264e-03	3.323909e-03	2.695691e+00
[16,]	15	2.789726e-03	2.428318e-03	1.969366e+00
[17,]	16	2.430096e-03	1.774033e-03	1.438741e+00
[18,]	17	2.154294e-03	1.296039e-03	1.051088e+00

Tabelle 1, Relative, erwartete und geometrisch verteilte Werte der simulierten Daten für den Chi-Quadrat-Test
Quelle: Eigene Darstellung

Die Sichtung der Daten ist deshalb erforderlich, weil alle Klassen [i] erwartete Werte von ≥ 5 haben sollen. Nach Ansicht von W.A. Stahel (2008) ist dieser Grenzwert zwar zu streng, da er aber sehr häufig in der Literatur verwendet wird, wird auch hier daran festgehalten.⁴⁷

Tabelle 1 zeigt deutlich, dass die Klassen 1-13 jeweils einen Wert > 5 aufweisen, womit keine Änderung nötig ist. Alle nachfolgenden Klassen haben allerdings Grenzwerte < 5 , weshalb sie

⁴⁶ Vgl. Hatzinger/Hornik/Nagel (2011), S. 266

⁴⁷ Vgl. Stahel (2008), S. 246

in neue Klassen zusammengefasst werden müssen, dessen Grenzwert dann über fünf liegen muss. Die Wahl der Zusammenlegung kann grundsätzlich frei getroffen werden. Analog zu W.A. Stahel (2008) werden daher zuerst Klasse 14 und 15 zusammengelegt. Rein praktisch werden also nicht nur die Erwartungswerte aufaddiert bis sie einen Grenzwert über 5 ergeben, sondern auch die dazugehörigen Werte der geometrischen Verteilung und der relativen Häufigkeiten. Da die übrigen Erwartungswerte sehr schnell recht klein werden kann man alle weiteren Klassen von 16 bis 811 zu einer Klasse zusammenlegen. Dessen Grenzwert liegt zwar knapp unter fünf, allerdings kann 1/5 der Werte auch unter fünf liegen.⁴⁸

Analog zu R. Hatzinger, K. Hornik, H. Nagel (2011) wird anschließend der Chi-Quadrat-Test über die neuen Klassen gelegt, was folgendes Ergebnis liefert:

```
X-squared
32.761924357
p.wert  0.003117558
```

Die Auswertung zeigt offensichtlich einen P-Wert der kleiner als das allgemein gültige Signifikanzniveau von 0,05 ist. Die Nullhypothese muss also verworfen werden, was bedeutet, dass es also doch signifikante Unterschiede zu den Erwartungen aus einer geometrischen Verteilung gibt. Offensichtlich kann also keine Basis für eine Gesetzmäßigkeit in der Struktur der Trends gefunden werden.

7 Die Trendstrukturanalyse des EURO STOXX 50

Wie weiter oben bereits angedeutet, besteht ein wesentlicher Teil dieser Arbeit auch darin, mit Hilfe des entwickelten Algorithmus echte historische Daten zu analysieren und mit den Ergebnissen der simulierten Daten zu vergleichen. Dabei wurde der EURO STOXX 50 gewählt. Dieser ist an sich der wichtigste Aktienindex in der Eurozone in dem die 50 größten börsennotierten Unternehmen aller Wirtschaftssektoren enthalten sind. Um sicher zu gehen, dass auch immer die größten 50 Unternehmen im EURO STOXX 50 notieren, wird vor allem jeweils die Marktkapitalisierung nach Streubesitz einmal jährlich überprüft. Allgemein gilt der EURO STOXX 50 als Indikator für die Entwicklung bzw. Situation am europäischen Aktienmarkt.⁴⁹ In einem Abschnitt weiter oben wurde bereits erwähnt, dass sich die konkreten historischen Daten über einen Zeitraum zwischen 2005 und 2013 erstrecken. Jedes Handelsjahr umfasst dabei

⁴⁸ Vgl. Stahel (2008), S. 246f.

⁴⁹ Vgl. ONVISTA, [Zugriff am 14.08.2015]

rund 260 Handelstage, wobei jeder Handelstag 50400 Ticks zählt. Da eine solche Fülle von Daten technisch schwieriger zu verarbeiten ist und recht unübersichtlich ist, werden die Vektoren welche die Ticks enthalten komprimiert. Geht man davon aus, dass ein Tick etwa einer Sekunde an einem Handelstag entspricht, macht es also eher Sinn über jeweils 60 Ticks einen Durchschnittswert zu berechnen. Somit entstehen pro Tag 840 Preisbewegungen die etwa einer Minute entsprechen.

Die technische Vorgangsweise zur Analyse der EURO STOXX 50 Daten wird weiter im Anhang in Abschnitt 3 genau erläutert. Im Wesentlichen wird genau derselbe Algorithmus, wie er bereits oben angewendet wurde, auch in diesem Fall verwendet um später einen Vergleich mit den simulierten Daten zu ermöglichen. Allerdings bestehen doch kleine Unterschiede im Code, wie dies im Anhang zu sehen ist. Anders als bei den simulierten Kursen, wird hier der eigentliche Algorithmus nicht mittels einer R-Funktion mehrmals wiederholt. Vielmehr werden alle Handelstage einzeln von 2005 bis 2013 zu Beginn mittels „Lapply“-Schleife in eine R-Liste gespeichert. Diese Liste hat somit genau so viele Einträge (d.h. Vektoren, bestehend aus den Werten der jeweiligen Handelstage) wie es Iterationen bei den simulierten Brownschen Bewegungen gibt.

Im weiteren Verlauf wendet die Lapply-Schleife den Algorithmus auf alle Vektoren der vorher genannten Liste an und erstellt analog zur Beschreibung weiter oben die Häufigkeitsverteilung der verschiedenen langen Trends. Das Ergebnis ist in der folgenden Abbildung zu sehen.

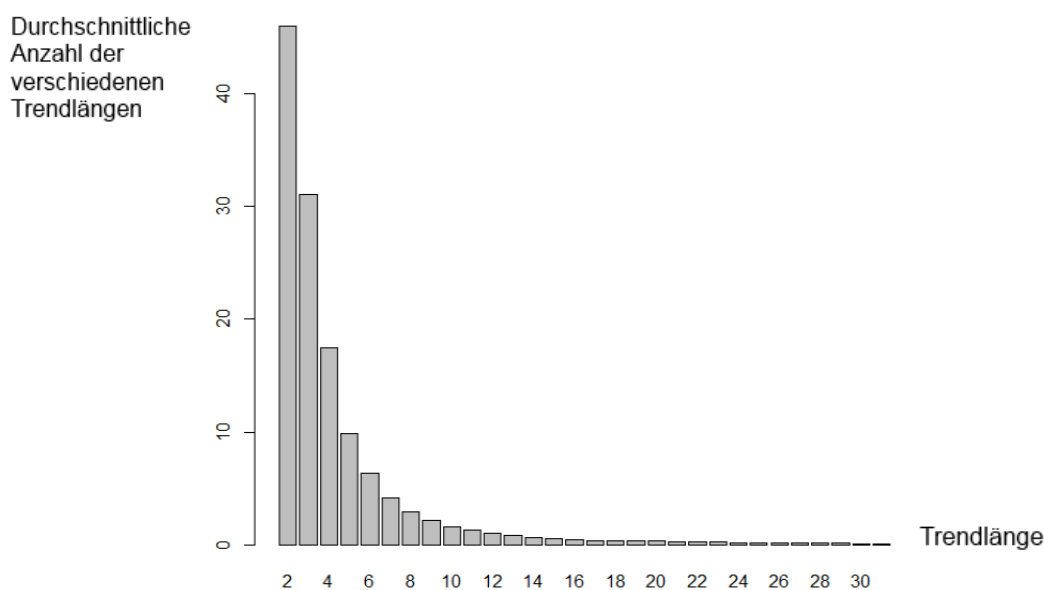


Abbildung 10, Die Trendstruktur der "Masterverteilung" der historischen Daten
Quelle: Eigene Darstellung

Abbildung 10 zeigt also, analog zur Anzahl der Iterationen bei den simulierten Daten, das Ergebnis aus 2131 vergangenen Handelstagen. Auch hier dominieren vor allem sehr kurze Trendverläufe. So kommen etwa 2er Trends im Schnitt rund 46-mal an einem Handelstag vor, während z.B. 3er bzw. 4er Trends 31-mal bzw. 17-mal vorkommen. Sehr ähnlich wie oben kommen auch hier etwas über 400 verschiedene Trendlängen unter allen Tagescharts vor. Einigermäßen verblüffend ist aber auch die Struktur der Verteilungskurve welche auch hier sehr schnell abflacht. Während bei den simulierten, geometrischen Brownschen Bewegungen Trendlängen bis 11 bzw. 12 im Durchschnitt mindestens einmal pro Chart vorkommen, sind es bei den realen Daten Trendlängen von 12 Punkten, welche mindestens einmal vorkommen. Auch hier sind zumindest zwischen 2 und 12 alle Trendlängen vorhanden. Lediglich am Ende der Verteilungen ist zu beobachten, dass bei den simulierten Daten am Ende eine Hand voll längere Trendverläufe vorhanden sind, die aber in allen 2131 Iterationen nur einmal oder zweimal vorkommen. Im Allgemeinen sind sich die zwei Häufigkeitsverteilungen also durchaus ähnlich. Bevor aber etwaige Unterschiede bzw. Ähnlichkeiten graphisch analysiert und statistisch getestet werden, ist noch die Frage zu klären ob die Häufigkeitsverteilung der historischen EURO STOXX 50 Daten einer theoretischen (geometrischen) Verteilung folgt oder, wie es bei den simulierten Daten der Fall war, eher nicht.

8 Statistische Auswertung der EURO STOXX 50 Daten

Wie oben, ist es auch hier von Interesse ob die Häufigkeitsverteilung in Abbildung 9 einer theoretischen Verteilung folgt, oder nicht. Weil das Ergebnis bzw. die Häufigkeitsverteilung sehr stark der Verteilung der simulierten Daten ähnelt kann auch, unter denselben Voraussetzungen wie oben, getestet werden ob die Trendverteilung der EURO STOXX 50 Daten einer geometrischen Verteilung entspricht. Das Signifikanzniveau liegt bei 0,05. Dabei gilt:

H_0 : Die zu testende Verteilung zeigt keine großen Unterschiede zu den Erwartungen einer geometrischen Verteilung.

H_1 : Die zu testende Verteilung zeigt signifikante Unterschiede zu den Erwartungen einer geometrischen Verteilung.

Auch hier wird der Chi-Quadrat-Anpassungstest verwendet. Analog zu den Beschreibungen weiter oben, werden die relativen bzw. erwarteten Häufigkeiten sowie die Werte der geometrischen Verteilung berechnet.

Es muss selbstverständlich wieder eine Sichtung der Daten erfolgen, um eine adäquate Klasseneinteilung durchzuführen. Einen Auszug der berechneten Werte zeigt Tabelle 2.

	absolut	relativ	geometrisch	erwartete
[1,]	0	3.462066e-01	2.417732e-01	9.961055e+01
[2,]	1	2.339323e-01	1.833189e-01	7.552739e+01
[3,]	2	1.318194e-01	1.389973e-01	5.726689e+01
[4,]	3	7.393383e-02	1.053915e-01	4.342129e+01
[5,]	4	4.782310e-02	7.991066e-02	3.292319e+01
[6,]	5	3.150611e-02	6.059040e-02	2.496325e+01
[7,]	6	2.214727e-02	4.594127e-02	1.892780e+01
[8,]	7	1.650792e-02	3.483390e-02	1.435157e+01
[9,]	8	1.217322e-02	2.641200e-02	1.088174e+01
[10,]	9	9.585128e-03	2.002629e-02	8.250830e+00
[11,]	10	7.845591e-03	1.518447e-02	6.256000e+00
[12,]	11	6.399513e-03	1.151327e-02	4.743467e+00
[13,]	12	5.066576e-03	8.729670e-03	3.596624e+00
[14,]	13	4.366519e-03	6.619070e-03	2.727057e+00
[15,]	14	3.708889e-03	5.018757e-03	2.067728e+00
[16,]	15	3.030046e-03	3.805356e-03	1.567807e+00
[17,]	16	2.948726e-03	2.885323e-03	1.188753e+00

Tabelle 2, Relative, erwartete und geometrisch verteilte Werte der historischen Daten für den Chi-Quadrat-Test
Quelle: Eigene Darstellung

Die Betrachtung ist hier wieder deshalb wichtig, um herauszufinden welche Klassen [i] erwartete Werte ≥ 5 haben. Offensichtlich können also die Klassen 1 bis 11 so belassen werden, wie sie sind. Die restlichen Klassen müssen (frei) so zusammen gelegt werden, sodass dessen Erwartungswert mindestens dem genannten Grenzwert entsprechen. Dementsprechend werden also die Klassen 12 und 13 bzw. die Klassen 14 bis 16 zusammengelegt. Da die erwarteten Werte ab Klasse 17 sehr schnell relativ klein werden können, ähnlich wie oben, alle Klassen von 17 bis 412 (d.h. bis zum Ende) zusammengefasst werden.

Auch hier fällt bereits auf, dass die Struktur der Verteilung sehr ähnlich mit jener der simulierten Daten ist, wo die Klassen ab einer Trendlänge von 16 bis zum Ende zusammengelegt werden mussten. Eine genauere Analyse dieser Unterschiede folgt jedoch erst weiter unten.

Da nun alle Klassen dem Chi-Quadrat-Anpassungstest gerecht verarbeitet wurden, kann der Test über die neuen Klassen gelegt werden. Dies liefert folgendes Ergebnis:

X-squared
29.892733325
p.wert 0.004880192

Offensichtlich fällt hier das Ergebnis wieder sehr ähnlich, wie bei der Auswertung der simulierten geometrischen Brownschen Bewegungen, aus. Der P-Wert ist zwar ein wenig größer als jener weiter oben, allerdings ist die vorliegende Verteilung dennoch signifikant verschieden von einer geometrischen Verteilung. Das bedeutet, dass die Nullhypothese verworfen werden muss.

Somit kann auch hier keine Basis für eine Gesetzmäßigkeit in der Entwicklung der Anzahl der verschiedenen Trends gefunden werden. Allerdings gilt es darauf hinzuweisen, dass hier nur Anpassungstests auf die geometrische Verteilung durchgeführt wurden. Die vorliegenden Strukturen der verschiedenen Trends passen zwar höchstwahrscheinlich nicht zu den meisten gängigen theoretischen Verteilungen, aber es ist dennoch nicht ausgeschlossen, dass positive Ergebnisse mit spezielleren Verteilungen zu erzielen sind. Ganz abgesehen davon schließen die Analysen nicht aus, dass aufgrund einer doch recht gleichmäßigen Kurve in den Trendstrukturen eine Gesetzmäßigkeit beispielsweise auf Basis einer einfachen geometrischen Reihe gefunden werden kann.

9 Simulierte Daten vs. Historische Daten

Weiter oben wurde bereits beschrieben, dass sich die Masterverteilungen der simulierten Brownschen Bewegungen und jene der EURO STOXX 50 Daten sehr ähnlich sind. Vergleicht man die beiden Abbildungen 8 und 9 kann man auf den ersten Blick keine Unterschiede feststellen. Ein Blick auf Abbildung 9 lässt nur Unterschiede auf der x-Achse erkennen.

Da die analysierten Daten des EURO STOXX von historischen Kursen stammen, könnte es durchaus sein, dass es im Gegensatz zu den simulierten Kursen z.B. weniger sehr kurzfristige Trends, aber dafür mehr mittelfristige Trends gibt. Außerdem könnte es sein, dass es in der Realität so gut wie keine Situationen gibt, in denen Trends über mehrere hundert Punkte verlaufen. Die Tabellen der echten bzw. simulierten Daten geben genauere Auskunft über die Unterschiede. Wie bereits oben erwähnt, ist die Anzahl der verschiedenen Trends sowohl bei den

simulierten Daten als auch bei den historischen Daten sehr ähnlich. Bei beiden Versionen dominieren die sehr kurzfristigen 2er bzw. 3er Trends, wobei der nächst längere Trend fast ausschließlich weniger oft vorkommt. Die folgenden zwei Tabellen zeigen, zur Veranschaulichung, Ausschnitte der Häufigkeitstabellen beider Versionen. Dargestellt werden jene Trendlängen die im Schnitt mindestens einmal pro Tageschart vorkommen:

Simulierte Daten

2	3	4	5	6	7
5.224026e+01	3.599625e+01	1.898029e+01	1.014172e+01	6.381980e+00	4.264664e+00
8	9	10	11	12	
2.908494e+00	2.097137e+00	1.593149e+00	1.222900e+00	9.549507e-01	

EURO STOXX 50 Daten

2	3	4	5	6	7
4.594979e+01	3.104833e+01	1.749554e+01	9.812764e+00	6.347255e+00	4.181605e+00
8	9	10	11	12	
2.939465e+00	2.190990e+00	1.615673e+00	1.272173e+00	1.041295e+00	

Offensichtlich besteht der größte Unterschied bei den beiden häufigsten Trends, nämlich den 2er und 3er Trends. Alle anderen Häufigkeiten der nachfolgenden Trendlängen sind sich bis zum Ende hin sehr ähnlich. So sieht man auch bei beiden Versionen, dass die Verteilungskurve gerade um die Trendlänge 12 bei „Eins“. Das bedeutet, dass alle anderen Trends im Durchschnitt weniger oft als einmal pro Chart vorkommen. Erstaunlicherweise gibt es auch bei den größten Trends keine großen Unterschiede. Bei den simulierten Daten gehen die beiden größten Trends über 761 bzw. 781 Punkte, während die größten Trends bei den historischen Daten über 767 bzw. 772 Punkte gehen. Alle diese Trendlängen kommen unter allen Iterationen bzw. Tagescharts genau einmal vor. Um einen besseren Überblick über den gesamten Verlauf beider Häufigkeitsverteilungen zu bekommen, wird im nächsten Schritt ein Diagramm erstellt. Allerdings macht es wegen der größeren Datenlage keinen Sinn zwei Balkendiagramme zu erstellen und diese übereinanderzulegen. Mit dieser Methode würde man wohl keine Unterschiede erkennen.

Daher werden die einzelnen Punkte der Häufigkeiten im Koordinatensystem einfach verbunden, was tatsächlich eine Kurve (in Linienform) ergibt. Legt man also die Kurve der simulierten Daten und jene der historischen Daten übereinander, kann man sehr leicht Unterschiede feststellen.

Durchschnittliche Anzahl
der verschiedenen
Trendlängen

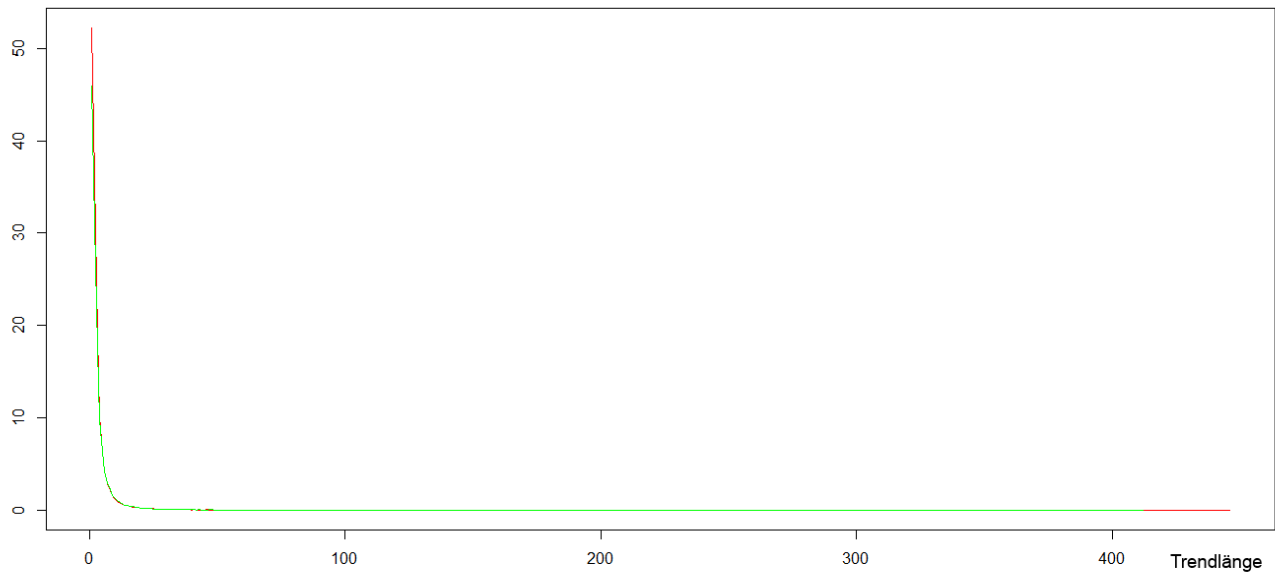


Abbildung 11, Liniendiagramm der übereinander gelegten Trendstrukturen der simulierten bzw. historischen Daten

Quelle: Eigene Darstellung

Abbildung 11 zeigt die Überlagerung beider Kurven. Dabei repräsentiert die rote Kurve jene der simulierten Masterverteilung, während die grüne Kurve die Daten der historischen EURO STOXX 50 Kurse repräsentiert. Es ist also auch hier deutlich zu erkennen, dass sich beide Verteilungen bloß zu Beginn ein wenig unterscheiden, aber ansonsten fast identisch sind.

Was anhand der Graphik in Abbildung 11 relativ deutlich zu erkennen ist, gilt es nun anhand eines statistischen Tests zu überprüfen. Dafür eignet sich der sogenannte Mann-Whitney-U-Test. Dieser ist ein nichtparametrischer Test zweier unabhängiger Stichproben, welcher überprüft ob die zentrale Tendenz von zwei verschiedenen Stichproben unterschiedlich ist.

Oft wird der Mann-Whitney-U-Test sonst angewendet um etwa Mittelwertunterschiede zwischen Experimental- und Kontrollgruppen zu ermitteln.⁵⁰

Die zentrale Frage für den Test lautet daher in diesem Fall:

Gibt es Unterschiede bezüglich der zentralen Tendenz zwischen der Verteilung der simulierten Daten und der historischen EURO STOXX 50 Daten?

⁵⁰ Vgl. Universität Zürich (2010), [Zugriff am 18.08.2015]

Diesbezüglich sehen die Nullhypothese und die Alternativhypothese folgendermaßen aus:

H_0 : *Die zwei zu testenden unabhängigen Stichproben unterscheiden sich in ihrer zentralen Tendenz nicht signifikant voneinander.*

H_1 : *Die zwei zu testenden unabhängigen Stichproben unterscheiden sich in ihrer zentralen Tendenz signifikant voneinander.*

Bei dem Mann-Whitney-U-Test wird das allgemein gültige Signifikanzniveau von 0.05 verwendet. Technisch wird also mit Hilfe des „`wilcox.test()`“ in R die zentrale Tendenz zwischen der Anzahl der Häufigkeiten der simulierten Daten und der historischen Daten überprüft, was zu folgendem Ergebnis führt:

data: anzahlsim and anzahlhist

W = 85128, p-value = 0.05981

Offensichtlich liegt der Grenzwert (p-value) knapp aber doch über dem Signifikanzniveau von 0,05. Mann kann also davon ausgehen, dass kein statistisch signifikanter Unterschied der zentralen Tendenz zwischen den Häufigkeitsverteilungen der simulierten Daten und der historischen EURO STOXX 50 Daten vorliegt. Da die Trendstruktur der simulierten Daten offensichtlich jener der historischen Daten sehr ähnlich ist, kann man behaupten, dass nach vorliegender Analyse langfristig wohl mehr Zufall in Preisbewegungen steckt als angenommen.

Dabei muss aber deutlich gemacht werden, dass im Rahmen dieser Analyse nur ein kleiner Teilbereich untersucht wurde. Zum einen wurde nur der EURO STOXX 50 zwischen 2005 und 2013 zum Vergleich mit den Simulationen herangezogen. Es ist nicht auszuschließen, dass ein Vergleich mit anderen historischen Daten ein anderes Ergebnis liefert. Außerdem wurde der Drift in den geometrischen Brownschen Bewegungen entfernt um keine künstlichen Trendverläufe einzubauen. Das Ergebnis war zwar signifikant, allerdings nur sehr knapp. Es ist nicht ausgeschlossen, dass kein signifikanter Zusammenhang zwischen simulierten und historischen Daten gefunden werden kann, wenn eine neue Masterverteilung, bestehend aus 2131 neuen Preisbewegungen, erstellt wird. Allerdings hat die Erfahrung gezeigt, dass simulierte Masterverteilungen, die aus wesentlich mehr simulierten Brownschen Bewegungen bestehen, im Schnitt auch nicht besonders verschieden sind. Natürlich kann auch der hier verwendete Algorithmus hinterfragt werden.

Grundsätzlich baut dieser auf den Ideen der „Technischen Analyse“ bzw. der Alliot-Wave-Theorie auf und unterscheidet bei der Erstellung von Trendlinien, ob sich beispielsweise ein Aufwärtstrend unter oder über einer steigenden Trendgeraden befindet. Dennoch kann man nicht hundertprozentig von Trendkanälen sprechen, wie sie sonst in der „Technischen Analyse“ verwendet werden. Trotzdem wurde der Algorithmus auf sämtliche historische und simulierte Preisbewegungen in dieser Arbeit angewendet, was etwaige Fehler in der Theorie des Algorithmus selbst möglicherweise relativiert.

10 Zusammenfassung

Im Zuge dieser Arbeit wird erörtert, welchen Stellenwert der Zufall an den Finanzmärkten hat. Derzeit ist in der Lehre der mathematisch-statistischen Analyse noch immer die Theorie der effizienten Märkte, weit verbreitet. Deren Grundprinzipien bzw. Voraussetzungen finden sich beispielsweise auch in Portfolioberechnungen wieder. Eines der wichtigsten Prinzipien ist dabei, dass wirtschaftliche Güterpreise nicht korreliert sind und einer Normalverteilung folgen. Doch Benoit Mandelbrot fand bereits im Zuge seiner Untersuchungen von Verteilungen bestimmter Arbeitseinkommen heraus, dass diese keiner Normalverteilung folgten und skaleninvariant waren, was gegen reine Zufallsprozesse spricht. Ähnliche Entdeckungen machte auch Harold Hurst als er Wasserstände am Nil untersuchte. Der nach ihm benannte „Hurst-Exponent“ beschreibt mit Werten von deutlich über 0,5 vermeintliche Zufallsbewegungen an den Finanzmärkten besser. Daraus lässt sich ableiten, dass Preisbewegungen wohl doch keine Zufallsprozesse sind. Trotzdem wird die Brownsche Bewegung noch immer benutzt um etwa Aktiencharts zu simulieren um dann beispielsweise Optionspreise zu berechnen. Obwohl Mandelbrot die Möglichkeit der Autokorrelation der Brownschen Bewegung hinzufügte und wir heute also die „fraktale Brownsche Bewegung“ kennen, besteht diese als stochastischer Prozess noch immer aus reinen Zufallsvariablen.

Vertreter der „Technischen Analyse“ sind hingegen entschiedene Gegner der Market-Efficiency-Hypothese. Sie sind davon überzeugt, dass sich massenpsychologisches Verhalten in immer wiederkehrenden Mustern in den Charts der Finanzmärkte finden lässt. Dabei stützen sich technische Analysten auf die Entdeckung von Trendverläufen in Preisbewegungen und versuchen mit ihrer Hilfe Wissen aus der Vergangenheit für die Zukunft zu nutzen. Allerdings lässt sich der Erfolg der „Technischen Analyse“ noch nicht wissenschaftlich festmachen. Denn obwohl bereits sehr viele kleinere Anleger aber auch institutionelle Anleger nach den Grundideen von Charles Dow an den Finanzmärkten handeln, konnte man die klare Überlegenheit der „Technischen Analyse“ gegenüber einer Buy-and-Hold-Strategie noch nicht beweisen. Dennoch gibt es öfter in der Literatur vielversprechende Hinweise zur Prognosequalität der „Technischen Analyse“, denen in Zukunft noch genauer nachgegangen werden muss.

Der selbst entwickelte Teil dieser Arbeit soll genau in die Diskussion über die Rolle des Zufalls in Aktiencharts einsteigen. Dabei erstellt ein Code nach einer bestimmten Logik zunächst in simulierten Bewegungen verschieden lange Trendgeraden, die auch öfter vorkommen können. Dieser Vorgang wird öfter wiederholt und Durchschnittswerte berechnet, die aussagen, wie oft

die verschiedenen Trendlängen durchschnittlich in einem Beobachtungszeitraum vorkommen. Die so entstandene „Masterverteilung“ wird hinsichtlich einer theoretischen Verteilung untersucht. Würde die Trendstruktur dieser Verteilung einer theoretischen Verteilung folgen, könnte man in weiterer Folge leichter ein Entwicklungsgesetz finden, mit dem man im Sinne der „Technischen Analyse“ statistisch die Wahrscheinlichkeit für eine bestimmte Trenddauer voraussagen könnte. Konkret kann aber eine geometrische Verteilung sowohl bei den simulierten Daten als auch bei den historischen Daten nicht nachgewiesen werden, was aber nicht bedeutet, dass die jeweiligen Masterverteilungen nicht einer etwas exotischeren Verteilung folgen. Abgesehen davon kann möglicherweise ein Entwicklungsgesetz etwa in Form einer geometrischen Reihe gesucht werden, was aber den Rahmen dieser Arbeit sprengen würde.

Interessant ist der Vergleich der Ergebnisse der simulierten bzw. historischen Daten. Die Häufigkeitsdiagramme der verschieden langen Trends, welche im Schnitt im Beobachtungszeitraum vorkommen, sind offensichtlich sehr ähnlich. Obwohl man annehmen könnte, dass der hier entwickelte Algorithmus etwa bei der Untersuchung der historischen Kurse im Durchschnitt eine andere Trendstruktur ergibt, ist sie dennoch kaum von jener der simulierten Preisbewegungen zu unterscheiden. Auch ein statistischer Test auf Gleichheit von zwei Stichproben bestätigt, dass sich die beiden Verteilungen in ihrer zentralen Tendenz nicht voneinander unterscheiden. Daraus könnte man womöglich schließen, dass sich die historischen Werte des EURO STOXX 50 nicht signifikant von den simulierten, geometrischen Brownschen Bewegungen unterscheidet, was bedeuten könnte, dass im Schnitt und auf längere Sicht mehr Zufall in Preisbewegungen steckt als gedacht. Allerdings ist zu erwähnen, dass die Brownsche Bewegung hinsichtlich ihres fehlenden Drifts manipuliert wurde, was andererseits nichts daran ändert, dass ihre Werte zufällig gebildet werden.

Somit gibt es bis heute keine klaren Resultate bei den Versuchen Preisbewegungen an den Finanzmärkten vollständig zu modellieren bzw. vorherzusagen. Insgesamt überwiegen aber die Hinweise, dass sich Aktiencharts wohl nicht hundertprozentig zufällig bewegen, wie es in der gängigen Portfoliotheorie gelehrt wird. Dies ist in einer derartig vernetzten Welt mit ähnlichen psychologischen Grundeinstellungen auch verständlich. Mathematisch-statistisch dürfte derzeit die fraktale Brownsche Bewegung mit Hilfe chaostheoretischer Überlegungen versuchen die von Thoma (1997) beschriebenen Zyklen bzw. skaleninvarianten Musterwiederholungen zu erklären. Auf der praktischen Seite ist die „Technische Analyse“ der Finanzmärkte unter sehr vielen Marktteilnehmern sehr populär, auch wenn deren Prognosequalität noch nicht eindeutig bewiesen werden konnte.

11 Quellenverzeichnis

11.1 Bücher

1. Hafner M., (2004), Klassifikation und Analyse finanzwirtschaftlicher Zeitreihen mit Hilfe von fraktalen Brownschen Bewegungen, Peter Lang / Europäischer Verlag der Wissenschaften, Frankfurt am Main
2. Hatzinger R., Hornik K., Nagel H., (2011), R – Einführung durch angewandte Statistik, Pearson Studium Verlag: München
3. Heussinger W.H., 2002, Elliott-Wave-Finanzmarktanalyse, Gabler-Verlag: Wiesbaden
4. Kohn W., Öztürk R., (2010), Statistik für Ökonomen, Springer-Verlag: Berlin
5. Murphy J. J., (2007), „Technische Analyse“ der Finanzmärkte, 4. Aufl., FinanzBuch Verlag GmbH: München
6. Spremann K., (2006), Portfoliomanagement, 3. Aufl., Oldenbourg Wissenschaftsverlag GmbH: München
7. Stahel W. A., (2008), Statistische Datenanalyse: Eine Einführung für Naturwissenschaftler, 5. Aufl., Friedr. Vieweg & Sohn Verlag / GWV Fachverlage GmbH: Wiesbaden
8. Thoma B., (1997), Chaostheorie, Wirtschaft und Börse: Das neue Paradigma in den Wirtschaftswissenschaften, 2. Aufl., Oldenbourg Verlag: München

11.2 Zeitschriften

1. Dorfleitner G., Klein C., (2002), Kursprognose mit Hilfe der „Technischen Analyse“ – Eine empirische Untersuchung, Swiss Society for Financial Market Research (pp. 497 – 521)

11.3 Online-Quellen

1. Blankenberger S., Vorberg D., (1998), Die Auswahl statistischer Tests und Maße, <http://www.students.uni-marburg.de/~Cal/Zeug/Entscheidungsbaum.pdf>, [Zugriff am 13.08.2015]
2. Griesing K., (2012), Geometrische Brownsche Bewegung und Brownsche Brücke, https://www.statistik.tu-dortmund.de/fileadmin/user_upload/Lehrstuehle/Ingenieur/Mueller/Lehre/Seminar/SS2012/GriesingP2.pdf, [Zugriff am 10.08.2015]
3. ONVISTA, Fakten und Wissenswertes zu EURO STOXX 50, <http://www.onvista.de/index/EURO-STOXX-50-Index-193736>, [Zugriff am 14.08.2015]
4. Universität Zürich, (2010), Mann-Whitney-U-Test, <http://www.methodenberatung.uzh.ch/datenanalyse/unterschiede/zentral/mann.html>, [Zugriff am 18.08.2015]

11.4 Dissertationen

1. Hofmann H. W., (1973), Empirische Untersuchung verschiedener Anlagestrategien der technischen Aktienanalyse anhand von 100 deutschen Standardaktien über einen Zeitraum von 5 Jahren (1967 – 1973) und Vergleich mit einer durch Simulation ermittelten Zufallsauswahl bzw. einer Kauf- und Haltestrategie, Universität München

12 Anhang

12.1 R-Code

Abschnitt 1: Simulierte Daten: Algorithmusdurchführung bzw. Untersuchung auf die Häufigkeit und Länge etwaiger Trendverläufe

1. Schritt: Simulation der Brown'schen Bewegung im Rahmen einer Funktion

```
myfunc <- function(){r <- 0
sigma <- 1
xbrown <- 100
N <- 840
Tbrown <- 1
Delta <- Tbrown/N
W <- numeric(N+1)
tbrown <- seq(0,Tbrown,length=N+1)
for(i in 2:(N+1)){W[i] <- W[i-1] + rnorm(1) * sqrt(Delta)}
S <- xbrown * exp((r-sigma^2/2)*tbrown + sigma*W)
x <- diff(S)
k <- NULL
```

2. Schritt: Erstellung der Leervektoren bzw. Leerlisten

```
d <- NULL
t <- vector("list",length(x))
A <- vector("list",length(x))
z <- vector("list",length(x)-2)
za <- vector("list",length(x)-2)
liste <- NULL
dreisum <- sapply(1:(length(x)-2), function(i) sum(x[c(i,(i+1))]))
```

3. Schritt: Auffüllen der Vektoren und Listen aus Schritt 2 mittels "sapply" bzw. "lapply"

```
dreisumi <- lapply(1:(length(x)-2), function(i) dreisum[i:(length(x)-2)])
zdreisumi <- lapply(1:(length(x)-4), function(i) dreisumi[[i]][3:length(dreisumi[[i]])]<0)
zadreisumi <- lapply(1:(length(S)-4), function(i) dreisumi[[i]][3:length(dreisumi[[i]])]>0)
Si <- lapply(1:(length(x)-2), function(i) S[i:(length(x))])
```

4. Schritt: Festlegen des Startwertes für die Iterationsvariable bzw. für die Dummy-Variable

```
i <- 1
h <- 1
```

5. Schritt: Eigentliche Durchführung des Algorithmus mittels "while"-Schleife

```
while(i<(length(x)-3) & h!=Inf){
k <- c(k,k <- (S[i]-S[i+2])/(-2))
d <- c(d,d <- (S[i+2]*i-S[i]*(i+2))/(-2))
t[[i]] <- i:(length(x))
A[[i]] <- k[length(liste)+1]*t[[i]]+d[length(liste)+1]
A[[i]][3] <- S[i+2]
z[[i]] <- Si[[i]][3:length(Si[[i]])]<A[[i]][3:length(A[[i]])]
za[[i]] <- Si[[i]][3:length(Si[[i]])]>A[[i]][3:length(A[[i]])]

if(k[length(liste)+1]>0 & S[i+3]>A[[i]][4] & is.element(TRUE,z[[i]])){h <-
(min(which(z[[i]]!=FALSE))+1)}else{
if(k[length(liste)+1]>0 & S[i+3]<A[[i]][4] & is.element(TRUE,za[[i]])){h <-
(min(which(za[[i]]!=FALSE))+1)}else{
if(k[length(liste)+1]<0 & S[i+3]>A[[i]][4] & is.element(TRUE,z[[i]])){h <-
(min(which(z[[i]]!=FALSE))+1)}else{
if(k[length(liste)+1]<0 & S[i+3]<A[[i]][4] & is.element(TRUE,za[[i]])){h <-
(min(which(za[[i]]!=FALSE))+1)}else{
if(k[length(liste)+1]>0 & S[i+3]>A[[i]][4] & (all(z[[i]]==FALSE))){h <-
(min(which(zdreisumi[[i]]!=FALSE))+2)}else{
if(k[length(liste)+1]>0 & S[i+3]<A[[i]][4] & (all(za[[i]]==FALSE))){h <-
(min(which(zdreisumi[[i]]!=FALSE))+2)}else{
if(k[length(liste)+1]<0 & S[i+3]>A[[i]][4] & (all(z[[i]]==FALSE))){h <-
(min(which(zadreisumi[[i]]!=FALSE))+2)}else{
if(k[length(liste)+1]<0 & S[i+3]<A[[i]][4] & (all(za[[i]]==FALSE))){h <-
(min(which(zadreisumi[[i]]!=FALSE))+2)}}}}}}}}}}

liste <- c(liste,i)
i <- i+h-1
if((length(x)-3)<=i & i<=length(x)){liste <- c(liste,i)}
return(liste)}
trendtest <- replicate(2131,myfunc(),simplify=FALSE)
```

6. Schritt: 2131-malige Wiederholung der gesamten Funktion (des gesamten Ablaufes Schritt 1-5)

```
trendtestdiff <- lapply(1:2131, function(i) diff(trendtest[[i]]))
```

7. Schritt: Erstellung eines Balkendiagramms

```
trendtestdiffbind <- unlist(trendtestdiff)
trendmittel <- table(trendtestdiffbind)/2131
barplot(trendmittel)
```

Abschnitt 2: Historische Daten: Algorithmusdurchführung bzw. Untersuchung auf die Häufigkeit und Länge etwaiger Trendverläufe

Schritt 1: Einlesen der historischen Daten in R und Transformation in eine Liste

```
setwd("C:/Users/Documents/alle Daten")
ldf <- list()
csvliste <- dir(pattern="*.csv")
for(i in 1:length(csvliste)){ldf[[i]] <- read.csv(csvliste[i]),2}
list <- lapply(ldf, function(x) unname(sapply(split(x,
ceiling(seq_along(x)/(length(x)/840))), mean)))

trendtest <- lapply(list, FUN = function(Sgrob) {

if(Sgrob[1]==Sgrob[2]){
S <- Sgrob[(min(which(Sgrob!=Sgrob[1]))):length(Sgrob)]}else{
S <- Sgrob}
x <- diff(S)
k <- NULL
d <- NULL
etc.
```

Der Rest verläuft analog zu Abschnitt 1.

Abschnitt 3: Chi-Quadrat Test für simulierte bzw. historische Daten

1. Schritt: Klassen erzeugen

```
anzahl <- tabulate(trendtestdiffbind)
anzahl <- anzahl[! anzahl %in% 0]
anzahl <- anzahl/10
tr <- sort(unique(trendtestdiffbind))
tr2 <- tr-2
n <- sum(anzahl)
mw <- sum(tr2*anzahl)/n
pgeom <- 1/mw
relhbeob <- anzahl/n
relhgeom <- dgeom(tr2, pgeom)
geomvert <- c(relhgeom[1:(length(relhgeom)-1)],1-sum(relhgeom[1:(length(relhgeom)-
1)])) # die Summe muss 1 betragen
erwartet <- geomvert*length(trendmittel)
tab <- c(trendmittel[1:length(trendmittel)])
cbind(absolut=tr2, relativ=relhbeob, geometrisch=geomvert, erwartete=erwartet)
```

2. Schritt: Klassen zusammenlegen

```
anzahlklassenger <- c(anzahl[1:4],sum(anzahl[5:6]),sum(anzahl[7:length(anzahl)]))
```

3. Schritt: Sichtung der Daten. Die relativen Häufigkeiten müssen über fünf liegen.

```
erwartetklassenger <-  
c(erwartet[1:4],sum(erwartet[5:6]),sum(erwartet[7:length(erwartet)]))  
geomvertklassenger <-  
c(geomvert[1:4],sum(geomvert[5:6]),sum(geomvert[7:length(geomvert)]))  
  
ct <- chisq.test(anzahlklassenger,p=geomvertklassenger)  
p.wert <- 1-pchisq(ct$statistic, df=(length(anzahlklassenger)-1))  
rbind(ct$statistic,p.wert)
```

4. Schritt: Auszug bzw. Beispiel der Testergebnisse

	absolut	relativ	geometrisch	erwartete
[1,]	0	3.608868e-01	2.694393e-01	2.185153e+02
[2,]	1	2.485370e-01	1.968418e-01	1.596387e+02
[3,]	2	1.307028e-01	1.438049e-01	1.166257e+02
[4,]	3	6.937840e-02	1.050582e-01	8.520218e+01
[5,]	4	4.395544e-02	7.675137e-02	6.224536e+01

```
erwartetklassenger <-  
c(erwartet[1:13],sum(erwartet[14:15]),sum(erwartet[16:length(erwartet)]))  
anzahlklassenger <- c(anzahl[1:13],sum(anzahl[14:15]),sum(anzahl[16:length(anzahl)]))  
geomvertklassenger <-  
c(geomvert[1:13],sum(geomvert[14:15]),sum(geomvert[16:length(geomvert)]))  
  
> rbind(ct$statistic,p.wert)
```

Abschnitt 4: Test auf Gleichheit der simulierten bzw. historischen Daten

1. Schritt: Laden des Codes aus Abschnitt 1

2. Schritt: Wilcoxon-Test

```
anzahlsim <- tabulate(trendtestdiffbindsim)  
anzahlsim <- anzahl[! Anzahl %in% 0]  
anzahlsim <- anzahl/(length(list))  
anzahl <- tabulate(trendtestdiffbind)  
anzahl <- anzahl[! Anzahl %in% 0]  
anzahl <- anzahl/(length(list))  
  
wilcox.test(anzahl, anzahlsim)
```

12.2 Abstract (Deutsch)

Die vorliegende Arbeit behandelt etwaige Trends im Chaos der Finanzmärkte. Zunächst liegt der theoretische Ausgangspunkt bei der Theorie der effizienten Märkte, die in der Literatur weit verbreitet ist. Diese Market-Efficiency-Hypothese kann auch einen Großteil des Marktgeschehens erklären, aber eben nicht alles. Das Argument ihrer Anhänger ist, dass Aktienkurse rein zufällig entstehen und daher normalverteilt sind. Kritiker stützen sich unter anderem auf die Erkenntnisse der Chaostheorie. Diese belegen mit Hilfe des „Hurst-Exponenten“, dass Aktienkurse nicht zufällig entstehen. Außerdem kann man in der Chaostheorie Trendverläufe und Zyklen nachweisen, die es bei reinem Zufall nicht geben dürfte. Genau so sehen das auch Vertreter der „Technischen Analyse“. Für sie existiert kein Zufall in den Preisbewegungen, vielmehr spiegelt sich in ihnen massenpsychologisches Verhalten wider. Daher versuchen Charttechniker immer wiederkehrende Muster zu erkennen, um Aktienkurse vorherzusagen. Obwohl die Methode weit verbreitet ist, konnte bisher nicht klar bewiesen oder widerlegt werden, dass sie funktioniert.

Im selbst erstellten Teil der Arbeit soll mittels simulierter geometrischer Brownscher Bewegungen sowie historischer Daten ermittelt werden, wie oft verschieden lange Trendverläufe in einem Tageschart durchschnittlich vorkommen. Könnte man nachweisen, dass etwa die Trendstrukturen einer geometrischen Verteilung folgen, könnte man in weiterer Folge versuchen ein Gesetz zur Vorhersage von Trendlängen zu formulieren. Die durchschnittlichen Trendstrukturverläufe der simulierten und historischen Daten weisen zwar ein starkes Muster auf, die folgen aber keiner theoretischen Verteilung. Vergleicht man die beiden unterschiedlichen Strukturen miteinander, kann eine signifikante Gleichheit nachgewiesen werden. Somit könnte mehr Zufall in Preisbewegungen stecken, als gedacht.

12.3 Abstract (English)

This thesis deals with possible trends within chaos on financial markets. We start with the classical Market-Efficiency-Hypothesis which is still widely used in literature. It can still explain much of what is happening on the markets, but not everything. The main argument of its supporters is that stock movements are pure random and therefore normally distributed. Critics of this theory focus on the findings of the chaos theory instead. Those findings prove that price movements on the markets are not a matter of random. In additions to that the chaos theory

proves the existence of trend movements and cycles which can not exist in a random world. That's exactly how technical analysts see it. They don't think that random price movements exist, but that stock price movements reflect a lot of mass-psychological behaviour instead. That's why technical analysts look for perseverative models in stock charts to be able to predict them. Though this method is widely used, it has never been fully proven or disproven.

In the self made part of this thesis I am trying to find out about how often different lengths of trends both in simulated geometric Brownian Motions and in actual historical data occur on average. If it was possible to prove that the trend structures of both kinds of data followed some geometrical distribution, it would be maybe possible to find a law to predict the length of trends statistically. Though the trend structures of both kinds of data show a perfect curve, they don't follow any theoretical distribution. Anyway, a very interesting thing happens, when the simulated trend structures are tested against the trend structures formed by historical data. The results show that they are statistically equal. So one could assume that there's still more fortune in stock price movements than we like to have it.

12.4 Curriculum Vitae

Persönliche Daten:

Name: Florian Göschl
Titel: Bakk.rer.soc.oec.
Nationalität: Österreich
Geburtsdatum: 03.06.1988

Ausbildung:

seit Oktober 2011: *Masterstudium der Betriebswissenschaften* am Betriebswirtschaftlichen Zentrum der *Universität Wien*, 1090 Wien

Ausbildungsschwerpunkt: *Controlling und Finanzdienstleistungen*

September 2007 – Juli 2011: *Bakkalaureat-Studium der Betriebswissenschaften* am Betriebswirtschaftlichen Zentrum der *Universität Wien*, 1090 Wien

September 1998 – Juli 2006: *Bundesrealgymnasium Hagenmüllergasse*, 1030 Wien

September 1994 – Juli 1998: *Volksschule Wittelsbachstrasse*, 1020 Wien

12.5 Eidesstattliche Erklärung

Ich erkläre hiermit an Eides Statt, dass ich die vorliegende Masterarbeit selbstständig und ohne Benutzung anderer als der angegebenen Hilfsmittel angefertigt habe.

Die aus fremden Quellen direkt oder indirekt übernommenen Gedanken sind als solche kenntlich gemacht.

Die Arbeit wurde bisher in gleicher oder ähnlicher Form keiner anderen Prüfungsbehörde vorgelegt und auch noch nicht veröffentlicht.

Wien, am 16.11.2015