



universität
wien

MASTERARBEIT / MASTER'S THESIS

Titel der Masterarbeit / Title of the Master's Thesis

„Naive Bayes Investing“

verfasst von / submitted by

Patricia Wöber, BSc (WU)

gemeinsam mit / together with

Maximilian Url, B.A. (Econ.)

angestrebter akademischer Grad / in partial fulfilment of the requirements for the degree of

Master of Science (MSc)

Wien, 2021 / Vienna, 2021

Studienkennzahl lt. Studienblatt /
degree programme code as it appears on
the student record sheet:

UA 066 915

Studienrichtung lt. Studienblatt /
degree programme as it appears on
the student record sheet:

Masterstudium Betriebswirtschaft UG 2002

Betreut von / Supervisor:

ao. Univ.-Prof. Mag. Dr. Christian Keber

Inhaltsverzeichnis

Tabellenverzeichnis	iii
Abbildungsverzeichnis	v
1 Einleitung und Problemstellung PW	1
2 Literatur und theoretischer Hintergrund PW	3
2.1 Markteffizienzhypothese PW	3
2.2 Anlagestrategien – Diskussion der Markteffizienzhypothese PW	7
3 Klassifizierung MU	9
3.1 Machine Learning und Methodik der Klassifizierung MU	9
3.2 Erläuterung und Vergleich unterschiedlicher Klassifikatoren MU	13
3.2.1 Entscheidungsbaum-Klassifikator MU	13
3.2.2 Kritische Würdigung des Entscheidungsbaum-Klassifikators MU	20
3.2.3 Regel-basierter Klassifikator MU	21
3.2.4 Kritische Würdigung des regel-basierten Klassifikators MU	27
3.2.5 Klassifikator der k-nächsten Nachbarn („kNN“) MU	28
3.2.6 Kritische Würdigung des Klassifikators der k-nächsten Nachbarn MU	31
3.2.7 Naive Bayes Klassifikator MU	32
3.2.8 Kritische Würdigung des Naive Bayes Klassifikators MU	37
3.3 Auswahl des anzuwendenden Klassifikators MU	38
4 Überführung einer Long List an makroökonomischen Indikatoren in eine Short List PW	38
4.1 Lineare Regressionen PW	39
4.2 Multiple lineare Regressionen PW	40
4.3 Herausforderungen in der statistischen Analyse von Zeitreihen PW	41
4.4 Augmented Dickey Fuller Test und Differenzbildung PW	42
4.5 Multikollinearität und VIF-Test PW	47
4.6 Schrittweise Selektion der Variablen PW	50
5 Auswahl und Beschreibung relevanter Daten PW	52
5.1 Makroökonomische Indikatoren und deren Einfluss auf Aktienmärkte PW	52
5.1.1 Definition und Klassifizierung PW	53
5.1.2 Bedeutung makroökonomischer Indikatoren für den Aktienmarkt PW	54
5.2 Auswahlprozess und Erläuterung volkswirtschaftlicher Indikatoren PW	57
5.2.1 Ökonomische Indikatoren mit potentieller Relevanz für die Vergleichsmärkte PW	57
5.2.2 DAX ® PW	59

5.2.3	S&P 500 ® PW	60
5.3	Charakteristika gewählter Benchmarks PW	60
5.3.1	DAX ® PW	61
5.3.2	S&P 500® PW	65
6	Erläuterung und formale Beschreibung der Handelsstrategien MU	68
6.1	Strategie 1 (Simple Strategie) PW.....	69
6.2	Strategie 2 (Durchschnitt für Entscheidung („DFE“), exkl. Short Positionierung) MU.	71
6.3	Strategie 3 (Durchschnitt für Entscheidung („DFE“), inkl. Short Positionierung) MU..	75
6.4	Strategie 4 (Historisch optimale Entscheidung („HOE“), keine Short Positionierung) PW	77
7	Empirische Auswertung der Daten	80
7.1	Überführung einer Long List in eine Short List PW.....	80
7.2	Anwendung des Naive Bayes Klassifikators MU	84
7.3	Anwendung der Handelsstrategien PW.....	88
7.3.1	Strategie 1 (Simple Strategie) PW	88
7.3.2	Strategie 2 (DFE, exkl. Short Positionierung) MU.....	91
7.3.3	Strategie 3 (DFE, inkl. Short Positionierung) MU.....	94
7.3.4	Strategie 4 (HOE, keine Short Positionierung) PW.....	95
7.4	Gesamtrendite und Vergleich mit der Benchmark MU.....	97
7.4.1	Deutscher Markt – DAX MU	97
7.4.2	Amerikanischer Markt – S&P 500 PW.....	101
7.4.3	Gegenüberstellung der Effektivität der Handelsstrategien in den Vergleichsmärkten PW.....	105
8	Limitationen MU	105
9	Zusammenfassung und Ausblick MU	108
	References	111
	Abstract (Deutsch)	115
	Abstract (English).....	116

Tabellenverzeichnis

Tabelle 1: Berechnungsbeispiel Gini-Index	18
Tabelle 2: Ermittlung der Metriken für $\geq$ Ungleichung	24
Tabelle 3: Ermittlung der Metriken für $\leq$ Ungleichung	25
Tabelle 4: Ermittlung der Metriken für verbleibende Beobachtungen	27
Tabelle 5: Beispielhafte Darstellung des kNN-Klassifikators	30
Tabelle 6: Ermittlung Mittelwert und Stichproben-Standardabweichung je Klasse	35
Tabelle 7: Ermittlung der Likelihoods sowie Multiplikation von Likelihoods und Prior (Zählergröße)	36
Tabelle 8: Ermittlung der a-posteriori Wahrscheinlichkeiten	37
Tabelle 9: Long List an makroökonomischen Indikatoren für den DAX	59
Tabelle 10: Long List an makroökonomischen Indikatoren für den S&P 500	60
Tabelle 11: Zusammensetzung des DAX	62
Tabelle 12: Top 10 der größten S&P 500 Unternehmen	66
Tabelle 13: Ableitung der Handelsentscheidung für Strategie 1 - Beispiel	71
Tabelle 14: Ermittlung der laufenden Durchschnitte für Strategie 2 - Beispiel	73
Tabelle 15: Ableitung der Handelsentscheidung für Strategie 2 - Beispiel	74
Tabelle 16: Ableitung der Handelsentscheidung für Strategie 3 - Beispiel	77
Tabelle 17: Ableitung der Handelsentscheidung für Strategie 4 - Beispiel	79
Tabelle 18: Ermittlung Mittelwert und Standardabweichung für Short List	85
Tabelle 19: Ermittlung der Likelihoods sowie Multiplikation von Likelihood und Prior (Zählergröße) für 2017	86
Tabelle 20: Ermittlung der a-posteriori Wahrscheinlichkeiten sowie Positionierung im Index für 2017	86
Tabelle 21: Ermittlung der erzielten Renditen der Benchmark sowie mithilfe der Strategie 1	87
Tabelle 22: Ermittlung der a-posteriori Wahrscheinlichkeiten sowie der Positionierung im Index für 2017 nach Strategie 1	89
Tabelle 23: Darstellung der erzielten Renditen der Benchmark und nach Strategie 1 für 2017	90
Tabelle 24: Ermittlung der DFE für 2017	91
Tabelle 25: Ermittlung der Handelsentscheidung nach Strategie 2 für 2017	92
Tabelle 26: Darstellung der erzielten Renditen der Benchmark und nach Strategie 2 für 2017	93
Tabelle 27: Ermittlung der Handelsentscheidung nach Strategie 3 für 2017	94
Tabelle 28: Darstellung der erzielten Renditen der Benchmark und nach Strategie 3 für 2017	95
Tabelle 29: Datentabelle, zur Ermittlung der kritischen Werte für 2017	96

Tabelle 30: Handelsempfehlung nach Strategie 4 für 2017	96
Tabelle 31: Darstellung der erzielten Renditen der Benchmark und nach Strategie 4 für 2017	97
Tabelle 32: Gegenüberstellung der Performance je Strategie und DAX	98
Tabelle 33: Jährliche Out- / Underperformance je Strategie - DAX.....	99
Tabelle 34: Gegenüberstellung der Performance je Strategie und S&P 500	102
Tabelle 35: Jährliche Out- / Underperformance je Strategie - S&P 500.....	102

Abbildungsverzeichnis

Abbildung 1: Binärer Entscheidungsbaum.....	14
Abbildung 2: Binärer Entscheidungsbaum mit den Attributen Handelsbilanz und Einzelhandelsumsätze	15
Abbildung 3: Binäre Entscheidungssituation	16
Abbildung 4: Aufsplittung anhand unterschiedlicher Attribute-Werte	17
Abbildung 5: Sektorale Zusammensetzung des DAX	63
Abbildung 6: Entwicklung des DAX 2000-2019	64
Abbildung 7: Sektorale Zusammensetzung des S&P 500	67
Abbildung 8: Entwicklung des S&P 500 von 2000-2019	68
Abbildung 9: Ableitung der Handelsentscheidung für Strategie 1	70
Abbildung 10: Ermittlung der Rendite für Strategie 1.....	71
Abbildung 11: Ableitung der Handelsentscheidung für Strategie 2	74
Abbildung 12: Ermittlung der Rendite für Strategie 2.....	75
Abbildung 13: Ableitung der Handelsentscheidung für Strategie 3	76
Abbildung 14: Ermittlung der Rendite für Strategie 3.....	77
Abbildung 15: Ableitung der Handelsentscheidung für Strategie 4	79
Abbildung 16: Ermittlung der Rendite für Strategie 4.....	80
Abbildung 17: Augmented Dickey-Fuller Test in R	82
Abbildung 18: Multiple lineare Regressionsgleichung in R	82
Abbildung 19: VIF-Test in R	83
Abbildung 20: Multiple lineare Regressionsgleichung nach Durchführung VIF-Test	83
Abbildung 21: Ergebnis der Forward Selection Methode.....	83
Abbildung 22: Ergebnis der Backward Elimination Methode	84
Abbildung 23: Jährliche Renditen im Vergleich - DAX	98
Abbildung 24: Entwicklung Portfoliowert je Strategie - DAX	100
Abbildung 25: Jährliche Renditen im Vergleich – S&P 500	101
Abbildung 26: Entwicklung Portfoliowert je Strategie – S&P 500.....	103

1 Einleitung und Problemstellung

Naive Bayes Klassifikatoren haben in den vergangenen Jahren aufgrund ihres engen Zusammenhangs mit dem so genannten „Machine Learning“ und ihrer einfachen Anwendbarkeit viel Aufmerksamkeit erlangt und finden in einer Vielzahl von Bereichen Verwendung. Neben alltäglichen Anwendungsgebieten wie der Filterung von unerwünschten Spam-Nachrichten kommt dieser Klassifikator auch bei der Entwicklung von Handelsstrategien zum Einsatz, beispielsweise in Verbindung mit der Analyse von technischen Indikatoren sowie online Finanznachrichten. Der Zusammenhang zwischen der Entwicklung makroökonomischer Variablen und Aktienkursen ist ein in der Literatur vielfach analysiertes Thema. So zeigen empirische Analysen etwa zum Teil einen deutlichen Einfluss von volkswirtschaftlichen Indikatoren auf die Entwicklung von Aktienindizes (Asprem, 1989). Andere Wissenschaftler wiederum kommen sogar zu der Schlussfolgerung, dass die Entwicklung von Aktienkursen neben dem – großteils durch psychologische Faktoren (etwa Herdenverhalten oder „Contrarian Behavior“) getriebenen – Investorenverhalten maßgeblich durch makroökonomische Einflüsse bedingt ist (Mende A., 1990).

Im Rahmen der vorliegenden Arbeit soll eine Investmentstrategie entwickelt werden, welche die Vorteile des Naive Bayes Klassifikators und den eben beschriebenen Zusammenhang zwischen makroökonomischen Indikatoren auf die Entwicklung von Aktienkursen kombiniert. Diese Arbeit soll Antwort darauf gegeben, ob mit der erarbeiteten Handelsstrategie über einen bestimmten Analysezeitraum eine Outperformance der gewählten Referenz-Aktienindizes erzielt werden kann. Die Handelsstrategie, die im Zuge der Masterarbeit entwickelt wurde, basiert auf einer Beobachtung bestimmter volkswirtschaftlicher Indikatoren. Welche volkswirtschaftlichen Indikatoren dabei Einfluss auf die Investitionsentscheidung haben, ist von der statistischen Analyse einer Grundgesamtheit an volkswirtschaftlichen Indikatoren abhängig. Nur jene Indikatoren, die relevante statistische Kriterien erfüllen, finden Eingang in die De-/Investitionsentscheidung. Aufgrund der im Vorfeld klar definierten Regeln hinsichtlich Kaufs- und Verkaufsentscheidungen sind diese unabhängig von jeglichen subjektiven Einschätzungen der Investoren. Neben der Objektivität der

Investitionsentscheidung zählen die freie Zugänglichkeit und simple Beschaffbarkeit der Indikatoren sowie die vergleichsweise einfache Anwendbarkeit der Strategie zu den wesentlichen Vorteilen von „Naive Bayes Investing“.

Zielsetzung der Arbeit

Im Rahmen der vorliegenden Masterarbeit soll eine Handelsstrategie entwickelt werden, die eine Outperformance im Vergleich zur gewählten Benchmark ermöglicht. Da die Handelsstrategie mit dem Ziel entwickelt wurde, dass diese möglichst intuitiv verständlich, objektiv und zudem mit relativ geringem Rechenaufwand verbunden ist, wurde der Naive Bayes Klassifikator herangezogen. Um die Anwendbarkeit der Strategie zusätzlich zu vereinfachen, sollten als Inputvariablen Daten dienen, die frei zugänglich sind. Aus dem genannten Grund wurden für die Analysen makroökonomische Indikatoren herangezogen.

Methodik

Für die Analyse werden Daten von Jänner 2000 bis Dezember 2019 herangezogen. Eine durch Literaturrecherche abgeleiteten Liste von volkswirtschaftlichen Indikatoren, die potenziell Einfluss auf die Marktentwicklung haben, wurde mittels statistischer Auswertung in eine verkürzte Liste an Indikatoren überführt. Die Überführung in die verkürzte Liste wird für jedes Jahr rollierend über einen rückwirkenden 4-jährigen Zeitraum neu ermittelt. Die volkswirtschaftlichen Indikatoren dienen als Inputfaktoren für den Naive Bayes Klassifikator, entsprechend diesem die Handelsentscheidung (kaufen, verkaufen, halten) getroffen wird.

Gliederung der Arbeit

In Kapitel 2 werden die drei Formen der Markteffizienzhypothese erläutert, eine Auswahl an bekannten Investmentstrategien beschrieben und auf die Motivation der Arbeit, eine möglichst simple, mit wenig Rechenaufwand verbundene Handelsstrategie zu entwickeln, übergeleitet. Kapitel 3 beschäftigt sich mit dem Thema der Klassifikation, einem Vergleich unterschiedlicher Klassifikatoren und der Begründung der Wahl des Naive Bayes Klassifikators für die weiteren Analysen. In Kapitel 4 wird die Methodik zur

Auswahl der Inputparameter für den Naive Bayes Klassifikator beschrieben (statistische Analysen). In Kapitel 5 wird basierend auf einer Literaturrecherche eine Liste an Indikatoren erstellt, die potenziell Einfluss auf die Entwicklung von Aktienmärkten haben. In Kapitel 6 werden die einzelnen Handelsstrategien basierend auf den a-posteriori Wahrscheinlichkeiten erläutert und formal vorgestellt. In Kapitel 7 findet schlussendlich die empirische Auswertung der Daten statt. Außerdem wird eine Übersicht über die erzielten Renditen gegeben. Die Arbeit schließt mit den Limitationen der durchgeführten Analysen sowie mit einer Zusammenfassung der gewonnen Erkenntnisse und einem Ausblick ab.

2 Literatur und theoretischer Hintergrund

In dem nachfolgenden Kapitel werden zunächst die Markteffizienzhypothese sowie ihre drei Ausprägungsformen erläutert. In weiterer Folge wird die Kernaussage dieser Hypothese, dass eine Outperformance des Marktes nicht möglich ist, durch die Beschreibung bekannter Investmentstrategien und deren Erfolge in Frage gestellt, wobei auch die mit der jeweiligen Handelsstrategie einhergehenden Problematiken dargelegt werden. In diesem Zusammenhang soll auf die Motivation dieser Arbeit übergeleitet werden. Neben einer simplen Anwendbarkeit und dem freien Zugang zu den erforderlichen Daten, war es Ziel, eine Investmentstrategie zu entwickeln, welche Problematiken bekannter Handelsstrategien eben nicht aufweist.

Wie bereits eingangs beschrieben, basiert die entwickelte Strategie auf der theoretischen Konzeption des Naive-Bayes Klassifikators, wobei die Klassifizierung und somit die Handelsempfehlung (kaufen, verkaufen, halten) auf Grundlage volkswirtschaftlicher Indikatoren abgegeben wird.

2.1 Markteffizienzhypothese

Die Markteffizienzhypothese (auch *Efficient Market Hypothesis*, in weiterer Folge daher als „EMH“ abgekürzt) ist eine der wichtigsten Paradigmen und Bedingungen der neoklassischen Kapitalmarkttheorie. In diesem Zusammenhang wird unter Kapitalmarkteffizienz die „Informations-Verarbeitungs-Effizienz“ verstanden (Kogan, 2009, S. 3). Neben den Arbeiten von Louis Bachelier und Paul A. Samuelson waren es

in erster Linie die Publikationen von Eugene Fama, welche die Bedeutung der Markteffizienzhypothese aufgezeigt haben: „The primary role of the capital market is allocation of ownership of the economy’s capital stock. In general terms, the idea is, a market in which prices provide accurate signals for resource allocation: that is, a market in which firms can make production-investment decisions, and investors can choose among the securities that represent ownership of firms’ activities under the assumption that security prices at any time “fully reflect” all available information. A market in which prices always „fully reflect“ available information is called “efficient”” (Fama, 1970, S. 383).

Im Gegensatz zur Theorie rationaler Entscheidung, die sich dem Entscheidungsprozess einzelner Marktindividuen widmet, stehen bei der EMH Entscheidungen auf aggregierter Ebene im Zentrum des Interesses. Der Marktpreis von Wertpapieren ergibt sich demnach aus dem Zusammenspiel von Angebot und Nachfrage. Kernaussage dieser Hypothese ist, dass die Preise von Wertpapieren in einem effizienten Markt sämtliche verfügbare Informationen berücksichtigen und dementsprechend einpreisen (Theis, 2014, S. 35f). In einem effizienten Kapitalmarkt entwickeln sich Aktienkurse daher so, als wären sämtliche wertrelevante Informationen der Gesamtheit an Investoren zugänglich. Die Analyse von Unternehmen zur Informationsgewinnung oder selbst Insider-Informationen können in effizienten Märkten daher nicht genutzt werden, um eine Outperformance zu erzielen. Somit ist „eine systematische gewinnbringende Strategie (...) in effizienten Märkten nicht existent“ (Kogan, 2009, S. 4).

Wie bereits erwähnt, nimmt Fama eine Untergliederung der EMH entsprechend der angenommenen Informationseffizienz vor: Abhängig von den jeweiligen Informationsarten, die der Aktienkurs bereits abbilden soll, wird zwischen der schwachen, halb- oder mittelstrengen sowie strengen Form der Markteffizienz unterschieden:

- Die *schwache Form* der Markteffizienz unterstellt, dass nur die Kurshistorie im Preis eines Wertpapiers eingepreist wurde. Mittels Einsatzes von technischen Analysen kann nach dieser Ausprägung der EMH keine Outperformance erzielt werden. „Die Technische Analyse basiert auf der Annahme, dass die zukünftige Entwicklung der Aktienpreise komplett aus dem Kursverlauf sowie teilweise aus der Umsatzvolumina der Vergangenheit bestimmt werden kann. Auf eine zusammenhängende

Begründung des Preisverlaufs wird dabei verzichtet“ (Kogan, 2009, S. 4). Der gezielte Einsatz von Insiderinformationen, aber auch Veröffentlichungen von Seite des jeweiligen Unternehmens oder Analysten, hingegen ermöglichen die Erwirtschaftung von Outperformance (Theis, 2014, S. 35f).

- Bei der *halbstrengen Form* der Markteffizienz hingegen wird unterstellt, dass sich nur jene Informationen im Marktpreis widerspiegeln, die der Öffentlichkeit zugänglich sind. Dazu zählen zum einen Informationen, die das Unternehmen selbst – freiwillig oder aufgrund regulatorischer Vorgaben – den Marktteilnehmern zur Verfügung stellt, zum anderen Daten, die etwa von Finanzanalysten erhoben werden (Theis, 2014, S. 35f). Bei dieser Variante ist die Erwirtschaftung einer Outperformance weder unter Einsatz einer technischen noch einer fundamentalen Analyse möglich. Entsprechend der Fundamentalanalyse hängt der zukünftige Unternehmenswert von der Substanz- und Ertragskraft eines Unternehmens ab. Ziel ist es, den fundamentalen Wert eines Unternehmens zu ermitteln (Kogan, 2009, S. 5). Dem fundamentalen Wert eines Unternehmens, der interne Planungen und Managementexpectations widerspiegelt, steht der aktuelle Marktwert, der einem Unternehmen an der Börse zugeschrieben wird, gegenüber. Dieser externe Wert reflektiert nicht die Erwartungen des Managements sondern vielmehr jene der Investoren. Weichen der interne und externe Unternehmenswert voneinander ab, kommt es also zu einer Über- oder Unterbewertung, so spricht man von einer Wertlücke. Entsprechend der halbstrengen Form der EMH existieren derartige Wertlücken allerdings nicht, wodurch sich ergibt, dass eine Fundamentalanalyse keinen Informationsgewinn stiften kann, auf dessen Basis Investoren eine Outperformance am Markt erzielen könnten (Theis, 2014, S. 19f). Unter Implikation der halbstrengen Form ist die Erwirtschaftung einer Outperformance lediglich dann möglich, sofern Marktteilnehmer über entsprechendes Insiderwissen verfügen (Theis, 2014, S. 35f).
- Bei der *strengen Form* der EMH wird unterstellt, dass alle Anleger stets Zugang zu sämtlichen verfügbaren Informationen haben und der Börsenpreis von Wertpapieren zu jedem Zeitpunkt diese bereits vollumfänglich berücksichtigt, so auch Insiderinformationen. Eine Outperformance kann daher keinesfalls erzielt werden

(Theis, 2014, S. 35f). „One would not expect such an extreme model to be an exact description of the world, and it is probably best viewed as benchmark against which the importance of deviations from market efficiency can be judged” (Fama, 1970, S. 414).

Voraussetzungen für die Gültigkeit der EMH sind ein Entscheidungsverhalten, das von Rationalität geprägt ist (Homo Oeconomicus), sowie die Abwesenheit von Informations- und Transaktionskosten (Theis, 2014, S. 36). Darüberhinausgehend müssen die Erwartungen der Anleger dahingehend homogen sein, ob gewisse Informationen eine wertsteigernde oder -senkende Wirkungsweise auf den Aktienkurs haben (Kogan, 2009, S. 5).

Aus der EMH ergibt sich somit direkt die Random-Walk-Hypothese bzw. Theorie der symmetrischen Irrfahrt, welche die Entwicklung von Aktienkursen im Zeitverlauf mathematisch beschreibt. „Diese Theorie besagt, dass der Kurs eines Finanztitels eine Zufallsvariable ist, die sich aus der Realisation des Kurses zum vorherigen Zeitpunkt und einem unsicheren Kurszuwachs zusammensetzt“ (Kogan, 2009, S. 6). Da sämtliche wertrelevanten Informationen im Zeitpunkt ihres Entstehens vollständig im Aktienkurs eingepreist werden, treten keine „Wertlücken“ auf. Der Kurs einer Aktie ändert sich demnach nur dann, wenn Informationen am Markt verfügbar sind, die für die Marktkapitalisierung eines Unternehmens von Relevanz sind. Da bewertungsrelevante Informationen nur zufällig auftreten, wirken die Preisverläufe von Aktien auf informationseffizienten Kapitalmärkten so, als ob sie einem Zufallspfad folgen würden (Breuer und Breuer, 2020).

Eng verknüpft mit der Theorie der effizienten Märkte ist zudem das „Informationsparadoxon“, das aus der Annahme der EMH folgt, dass kein Investor in der Lage ist, durch Analysen jeglicher Art oder selbst Insiderinformationen eine Outperformance zu erwirtschaften. Es bleibt demnach kein Anreiz für Investoren, derartige Mühen auf sich zu nehmen. Daraus ergibt sich schlussendlich die Problemstellung, wie überhaupt Preisänderungen von Aktien entstehen, wenn Anleger keinerlei Anreiz haben, Informationen zu erheben, entsprechend dieser zu

(de-) investieren und somit schlussendlich den Kursverlauf einer Aktie mitzubestimmen (Kogan, 2009, S. 6).

In der Realität ist es allerdings so, dass einzelnen Akteure am Markt ein- und derselben Aktie unterschiedliche Werte beimessen, da diese divergierende Präferenzen hinsichtlich Risiko oder Erwartungen an das fragliche Unternehmen aufweisen. Es findet daher ein täglicher Handel mit börsennotierten Unternehmensanteilen statt. Während der subjektive Wert eines Käufers über dem Marktwert liegt, zu dem die Aktie im jeweiligen Zeitpunkt gehandelt wird, liegt jener des Verkäufers unter dem momentanen Börsenpreis (Coenenberg et al., 2015, S. 39).

Viele Investoren, insbesondere institutionelle Anleger und Hedge-Fonds, treffen ihre Handelsentscheidungen gemäß Anlagestrategien. Der nachfolgende Abschnitt bietet einen kurzen Überblick prominenter Strategien und stellt somit die Gültigkeit der Markteffizienzstrategie infrage.

2.2 Anlagestrategien – Diskussion der Markteffizienzhypothese

„Als Anlagestrategie wird die Vorgehensweise bezeichnet, nach der private und institutionelle Anleger bei ihren Investmententscheidungen operieren. Die konjunkturelle Gesamtlage, das finanzielle Vermögen und die individuelle Risikobereitschaft haben Einfluss auf die Anlagestrategie“ (börsenNEWS.de, 2021). Anlagestrategien können abhängig von der Risikobereitschaft in konservative und aggressive Strategien untergliedert werden: Während aggressive Strategien auf einen möglichst hohen Gewinn abzielen und eine kurze Haltedauer typisch ist, zeichnen sich konservative Strategien durch längerfristige Investments in sichere dafür weniger rentable Finanzinstrumente sowie Risikodiversifikation aus (börsenNEWS.de, 2021).

Beschränkt man die Finanzinstrumente auf Aktien, so bietet eine weitere Klassifikationsmöglichkeit die zugrundeliegende Analyseform (Fundamentalanalyse, technische Analyse).

Die Fundamentalanalyse, die die Basis für das so genannte „Value Investing“ bildet und maßgeblich von Benjamin Graham geprägt wurde, gliedert sich in drei Teilschritte. Zunächst erfolgt eine Vorselektion, in deren Rahmen Unternehmen identifiziert werden,

die möglicherweise vom Markt unterbewertet werden und deren Aktien daher interessante Investitionsobjekte darstellen. Die Vorselektion ist erforderlich, da aufgrund zeitlicher und finanzieller Restriktionen nicht sämtliche börsengehandelten Unternehmen analysiert werden können. Die Vorselektion erfolgt entweder anhand einer konträren oder kennzahlenbasierten Suchstrategie. Während die konträre Strategie auf qualitativen Kriterien fußt, stützt sich die kennzahlenbasierte Strategie auf quantitative Faktoren. Im nächsten Schritt wird der innere Unternehmenswert bestimmt. Aufgrund emotionaler, spekulativer und irrationaler Handelsentscheidungen anderer Marktteilnehmer weicht der innere Wert häufig von jenem Wert ab, mit dem Unternehmensanteile an der Börse gehandelt werden (Venitz, 2019, S. 14-16). Anhand der ermittelten Diskrepanzen werden schließlich Investitions- beziehungsweise Desinvestitionsentscheidungen getroffen: Zeigt die Analyse zum Beispiel, dass ein Unternehmen an der Börse unterbewertet ist, wird ein steigender Aktienkurs erwartet, da stets davon ausgegangen wird, dass Wertdiskrepanzen lediglich temporär sind (Gothein, 1995, S. 7f). Es sind somit zwei Annahmen für die Fundamentalanalyse und das „Value Investing“ wesentlich: Zum einen wird unterstellt, dass der Preis einer Aktie und deren Wert Divergenzen aufweisen können, zum anderen, dass zumindest langfristig eine Annäherung an den wahren Wert stattfindet (Venitz, 2019, S. 14). Meist zeigen sich allerdings Anlagestrategien, die rein auf fundamentalen Unternehmensdaten aufbauen, wenig erfolgreich. Ein weiterer Nachteil ist, dass es insbesondere der konträren Vorselektion einer intersubjektiven Validierung mangelt – diese Suchprozesse führen so zu keinem eindeutigen oder überprüfbaren Ergebnis. Zudem beschränkt sich der qualitative Suchprozess auf eine recht geringe Anzahl an Unternehmen, wobei mit steigender Intensität, mit der Unternehmen analysiert werden, die Anzahl an Analyseobjekten sinkt (Venitz, 2019, S. 20f). Beim kennzahlenbasierten Selektionsprozess hingegen ist die Auswahl objektiv nachvollziehbar, gleichzeitig kann auf ressourcenschonende Art und Weise auch eine Vielzahl an Unternehmen analysiert werden. Hier ergeben sich Kritikpunkte daraus, dass der kennzahlenbasierte Selektionsprozess rein mechanisch ist und sich die Ergebnisverwendung oft als unreflektiert erweist. Zudem sind die Variablen, die für die quantitative Voranalyse herangezogen werden, anfällig für Verzerrungen durch bilanzpolitische Maßnahmen, wodurch sich eine mangelnde Vergleichbarkeit der

analysierten Unternehmen ergeben kann. Weiters können Inkonsistenzen entstehen, wenn der Berechnung der Eingangsvariablen unterschiedliche Vorgehensweisen zugrunde gelegt werden (Venitz, 2019, S. 29).

Bei technischen Analysen (etwa gleitender Durchschnitt, exponentielle Glättung, Chartformationen) wird versucht, mittels der Untersuchung von Kursverläufen die Entwicklung einzelner Unternehmen aber auch eines gesamten Marktes vorherzusagen und so eine Outperformance zu erzielen. Problematisch zeigt sich hier die fehlende theoretische Fundierung. Außerdem ist der praktische Erfolg derartiger Anlagestrategien umstritten (Gothein, 1995, S. 7).

An die eingangs erwähnte Anlagestrategie des „Naive Bayes Investing“ soll in den nächsten Kapiteln der Arbeit Schritt für Schritt – beginnend mit der Thematik Klassifizierung – herangeführt werden.

3 Klassifizierung

Dieses Kapitel der Arbeit widmet sich dem Thema der Klassifizierung. Es wird erarbeitet, worum es sich dabei handelt und für welche Zwecke die Klassifizierung eingesetzt wird. Im Weiteren werden unterschiedliche Methoden der Klassifizierung erläutert und analysiert. Auf Basis der durchgeführten Analysen wird die Entscheidung getroffen, welcher Klassifikator für die Generierung der Handelssignale zum Einsatz kommt. Dabei soll der Intention der Arbeit, eine einfache, mit objektiven Entscheidungen verbundene, ohne viel rechnerischem Aufwand umsetzbare Anlagestrategie zu finden, Rechnung getragen werden. Schlussendlich wird zugunsten des Naive Bayes Klassifikators entschieden.

3.1 Machine Learning und Methodik der Klassifizierung

Machine Learning ist ein Teilbereich der künstlichen Intelligenz, bei dem es darum geht, Computerprogramme zu schreiben, die es der „Maschine“ ermöglichen, Aufgaben zu erledigen, wie beispielsweise Vorhersagen zu treffen. Die Maschine soll dabei über ein erlerntes Modell aus den erhaltenen Inputs ein gewünschtes Ergebnis hervorbringen. Machine Learning kann auf Basis der, der Maschine bereitgestellten Informationen in mehrere Teilbereiche gegliedert werden, von denen in der Folge drei Arten kurz erläutert

werden. Um „Supervised Learning“ handelt es sich, wenn der Maschine „beschriftete“ Inputdaten im Lernprozess zur Verfügung gestellt werden. In der vorliegenden Arbeit handelt es sich bei diesen beschrifteten Inputdaten beispielsweise um die Veränderung der Einzelhandelsumsätze einer Periode, welche mit einer korrespondierenden Klasse (in der vorliegenden Arbeit beispielsweise „Aktienindex steigt“) beschriftet werden. Die Maschine bekommt die Aufgabe eine Funktion zu finden, um die Klasse eines Inputdatenpunkts zu erklären. Unter Supervised Learning fällt beispielsweise der in dieser Arbeit behandelte Sachverhalt der Klassifizierung. Um „Unsupervised Learning“ handelt es sich, wenn genau das Gegenteil im Lernprozess geschieht. Der Maschine werden folglich keine klassifizierten Daten zur Verfügung gestellt. Zielsetzung bei dieser Art des Lernens ist es, versteckte Strukturen in den Inputdaten zu erkennen, um diese nutzen zu können. Bei „Semi-Supervised Learning“ handelt es sich, wie der Name erahnen lässt, um eine Kombination der beiden zuvor genannten Methoden. Das Ziel dieser Art des Lernens ist verbesserte Vorhersagen zu erhalten als es mit reinem Supervised Learning möglich wäre (Mohammed et al., 2017, S. 4-11).

In der vorliegenden Arbeit wird auf den Bereich des Supervised Learnings und im Speziellen auf die Klassifizierung eingegangen.

Klassifizierung beschreibt den Tatbestand, durch welchen beobachtbare Daten mithilfe einer Funktion, dem sogenannten Klassifikator (beispielsweise Entscheidungsbäume, Naive Bayes Klassifikator), mit im Vorhinein festgelegten Klassen beschriftet werden. Grundlage für die Klassifizierung bilden historische Beobachtungen, die aus einem Attribute-Set sowie aus einer Klassen-Variablen bestehen, welche die Klasse der Beobachtung festlegt (Kantardzic, 2020, S. 198).

Das Attribute-Set kann als n-dimensionaler Attribute-Vektor $X = (x_1, x_2, \dots, x_n)$ verstanden werden, wobei die x_i konkrete Ausprägungen der Attribute $A_1, A_2, \dots, A_n$ darstellen (z.B. Veränderung der Einzelhandelsumsätze einer Periode um 3 %). Das Attribute-Set kann somit aus nur einem Attribut bestehen oder aber auch mehrere Attribute beinhalten – beispielsweise die Attribute „Manufacturing PMI“, „Veränderung der Einzelhandelsumsätze“ und „Handelsbilanz“. Die Attribute im Attribute-Set können

sowohl kontinuierliche Werte als auch diskrete Werte annehmen. Die Klassen-Variable kann hingegen nur diskrete oder kategorische Ausprägungen – beispielsweise „steigender Aktienkurs“ – aufweisen (Cichosz, 2015, S. 9; Han et al., 2012, S. 328). Jeder Klassifikator generiert dabei auf seine bestimmte Art und Weise ein Modell, welches die Beziehung zwischen den Attributen (mit deren konkreten Ausprägungen) im Attribute-Set und der Klassen-Variablen möglichst gut erklärt. Dabei ist es essentiell, dass das mithilfe der Input-Daten erlernte Modell in der Lage ist, auch unbekannte Datensätze möglichst fehlerfrei der entsprechenden Klasse zuzuordnen (Cichosz, 2015, S. 9-12).

Der Tatbestand der Klassifizierung lässt sich in Anlehnung an Tan et al. (2006, S. 146) wie folgt darstellen:

$$\text{Attribute Set } X (\text{Input}) \rightarrow \text{Klassifikator} \rightarrow \text{Klasse (Output)} \quad (1)$$

Die auf diese Weise erarbeiteten Klassifikatoren können für verschiedene Anwendungsbereiche genutzt werden, so zum Beispiel auch, wie in der vorliegenden Arbeit, um Klassen für bestimmte Attribute-Kombinationen vorherzusagen. Hierunter versteht man ein „Predictive Model“. Zusätzlich kann ein Klassifikator als ein „Descriptive Model“ fungieren, indem er die notwendigen Ausprägungen hervorhebt, welche von den Attributen angenommen werden müssten, um beispielsweise die Klasse „steigender Aktienindex“ zu erhalten (Tan et al., 2020, S. 136).

Klassifikatoren werden in der Praxis in einer Vielzahl von Sektoren und Anwendungsbereichen eingesetzt, so unter anderem auch im Medizinbereich (Shafaf and Malek, 2019), im Bereich der Landwirtschaft (Camargo and Smith, 2009) oder der Astronomie (Freed and Lee, 2013).

Shafaf und Malek (2019) verdeutlichen in ihrem Paper, welche Klassifikatoren in der jüngeren Vergangenheit im Besonderen in Zusammenhang mit der verbesserten Erkennung von Krankheiten, der Vorhersage von Aufnahme und Entlassung von Patienten sowie im Zusammenhang mit Triage Situationen besonders häufig angewandt wurden. Darunter finden sich beispielsweise bekannte Klassifikatoren wie

Entscheidungsbäume, der Naive Bayes Klassifikator aber auch komplexere Klassifikatoren wie Support Vector Machines.

In der Landwirtschaft finden Klassifikatoren unter anderem Anwendung in der Erkennung von Krankheitserregern in Pflanzen (Camargo and Smith, 2009). Camargo und Smith (2009) versuchen mithilfe von Bildern von Pflanzenblättern, die typische Krankheitsanzeichen aufweisen und einem Support Vector Machine-Klassifikator, die Erkennung von Krankheiten in der Pflanzenzucht zu verbessern. Eine Anwendung dieser Methode könnte laut Autoren vor allem in abgelegenen Regionen mit erschwerter Zugänglichkeit nützlich sein.

In der Astronomie (Freed and Lee, 2013) werden Klassifikatoren beispielsweise herangezogen, um Galaxien anhand deren Struktur einer bestimmten Art von Galaxien zuzuordnen. Aufgrund der Vielzahl von Objekten ist die manuelle Klassifikation generell nicht umsetzbar. Aus diesem Grund erfreut man sich ebenfalls der Anwendung von Klassifikationsverfahren unter Rückgriff auf, unter anderem, lichttechnische Daten, die die Struktur der Galaxie beschreiben. Im Paper von Freed and Lee (2013) wird dafür abermals ein Support Vector Machine-Klassifikator herangezogen, um die Galaxien basierend auf deren Struktur einzuordnen.

Nicht nur in den oben beispielhaft genannten Industrien, sondern auch im Bereich der Finanzwirtschaft finden Klassifikatoren immer häufiger Anwendung. Diese Anwendungen greifen beispielsweise auf die Klassifikation von Texten zurück (Chan and Chong, 2017). In Chan und Chong (2017) versuchen die Autoren mithilfe von Klassifikatoren Informationen aus Textbausteinen in Medien wie beispielsweise Nachrichtenseiten, Blogs oder Tweets, die unter Umständen essentielle Bedeutung für finanzielle Entscheidungen haben, zu erlangen, um diese zur Stimmungserkennung an den Finanzmärkten zu nutzen. Zum Einsatz kommt eine Ensemble Technik, bei welcher ein Entscheidungsbaum-Klassifikator sowie ein Support Vector Machine-Klassifikator kombiniert werden.

Gerlein et al. (2016) hingegen versuchen mit ihrer Arbeit Erkenntnisse für die Formulierung von Handelsentscheidungen basierend auf saisonalen Inputdaten sowie

zeitverzögerten Kursdaten und technischen Indikatoren zu liefern, die sich aus historischen Kursdaten errechnen lassen. Dabei versuchen sie im FOREX Markt, unter Zuhilfenahme der genannten Inputdaten, anhand einfacher Klassifikatoren profitable Handelsentscheidungen zu treffen. Die Inputs für die Klassifikatoren bilden neben Schlusskursen, Tageszeiten, Wochentagen, prozentualen Preisveränderungen auch technische Indikatoren wie der Relative Stärke Index oder der Williams %R Oszillator. Die Autoren berücksichtigen in ihrer Arbeit für ihre Analysen unter anderem den Naive Bayes Klassifikator, den Entscheidungsbaum-Klassifikator oder aber auch den regelbasierten Klassifikator. Diese Klassifikatoren sind geprägt durch einen geringen rechnerischen Aufwand in der Konstruktion sowie der Anwendung zur Generierung der Handelsstrategie. Die Performance einzelner Klassifikatoren konnte mithilfe verschiedener Anpassungen wie beispielsweise einer rollierenden Trainingsphase, um aktuelle Marktgegebenheiten zu berücksichtigen, verbessert werden.

3.2 Erläuterung und Vergleich unterschiedlicher Klassifikatoren

Dieser Abschnitt beschäftigt sich mit einem Vergleich unterschiedlicher Verfahren zur Klassifikation von Daten. Die Verfahren werden theoretisch vorgestellt und hinsichtlich Anwendbarkeit auf die Problemstellung der vorliegenden Arbeit analysiert und bewertet. Hierfür werden die populären Klassifikatoren Entscheidungsbäume, regel-basierte Klassifikatoren, der Klassifikator basierend auf den k-nächsten Nachbarn („kNN“) und der Naive Bayes Klassifikator herangezogen.

3.2.1 Entscheidungsbaum-Klassifikator

Ein Entscheidungsbaum kann als eine Aneinanderreihung unterschiedlicher Entscheidungssituationen verstanden werden, welcher in mehreren Ebenen aufgebaut ist. In jeder Ebene gibt es einen oder mehrere Knoten, bei welchen Entscheidungen getroffen werden (beispielsweise Aufsplittung der Beobachtungen oder Festlegen der Klasse). Auf oberster Ebene findet man den sogenannten „Wurzelknoten“. Dieser Knoten hat keinen Vorgänger und keinen oder mehrere Nachfolger. Als zweite Art von Knoten lassen sich sogenannte „innere Knoten“ identifizieren. Diese Art von Knoten hat jeweils einen Vorgänger und zwei oder mehr Nachfolger. Das Besondere an Wurzelknoten und inneren Knoten ist, dass bei diesen Knoten immer eine Entscheidung getroffen wird,

anhand derer die Beobachtungen in disjunkte Mengen getrennt werden. Dies wird so lange fortgesetzt bis man zu einem sogenannten „Blattknoten“ gelangt. Diese dritte und letzte Art von Knoten zeichnet sich dadurch aus, dass sie jeweils einen Vorgänger hat und keinen Nachfolger. Kennzeichen dieser Knoten ist, dass jedem eine Klasse zugeordnet wird. Folglich werden die Beobachtungen in einem Blattknoten alle derselben Klasse zugeordnet (Rastogi and Shim, 2000, S. 316f).

In Abbildung 1 ist ein beispielhafter binärer Entscheidungsbaum in Anlehnung an Tan et al. (2020, S. 141) dargestellt bevor in weiterer Folge auf die allgemeine Konstruktion von Entscheidungsbäumen eingegangen wird.

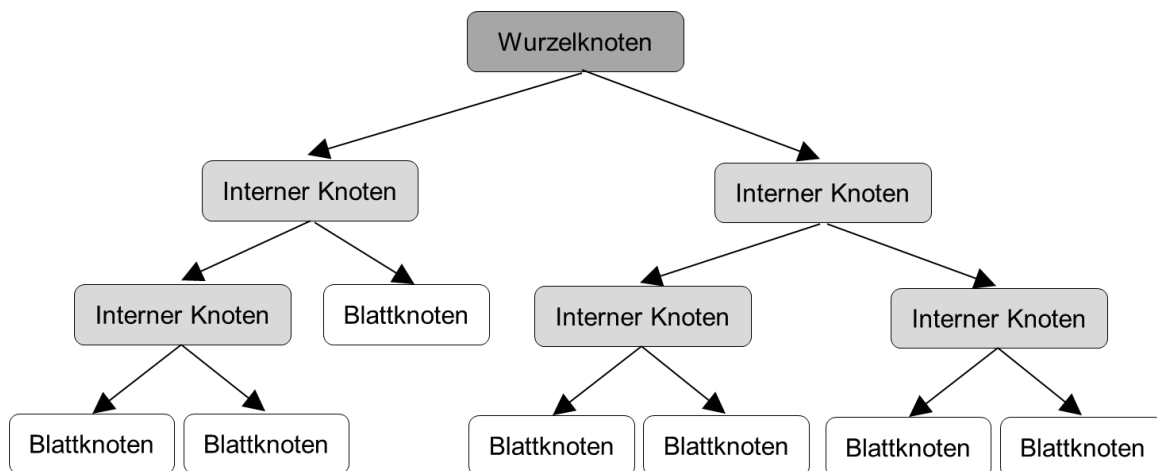


Abbildung 1: Binärer Entscheidungsbaum

Eine beliebte Methode, um Entscheidungsbäume zu konstruieren, stellt der „Hunt“-Algorithmus dar. Dabei wird der Entscheidungsbaum beginnend mit dem Wurzelknoten rekursiv konstruiert. Mithilfe eines Entscheidungskriteriums (beispielsweise anhand der Ausprägung des Attributs „monatliche Veränderung der Einzelhandelsumsätze“ zu einem bestimmten Zeitpunkt) werden die Beobachtungen im Knoten in disjunkte Teilmengen gespalten. Ziel davon ist es disjunkte Teilmengen zu erhalten, die möglichst nur einer Klasse zugeordnet werden können. Die Beobachtungen, die sich auf jedem Ast des Entscheidungsbaumes nach der Aufspaltung befinden, werden neuerlich mithilfe eines Entscheidungskriteriums aufgespalten. Der Algorithmus trennt die zur Verfügung stehenden Beobachtungen so lange weiter auf und fügt neue Knoten hinzu, bis entweder alle Beobachtungen in einem Knoten dieselbe Klasse aufweisen oder aber die

Beobachtungen alle dieselben Attribute-Werte aufweisen und dementsprechend nicht weiter aufgetrennt werden können. Schlussendlich bleiben am Ende in jedem Blattknoten nur Beobachtungen derselben Klasse übrig (Tan et al., 2020, S. 141-143).

Des Weiteren ist es möglich, den Entscheidungsbaumwachstumsprozess zu stoppen, wenn beispielsweise die Anzahl an Beobachtungen nach einem weiteren Aufspalten in einem inneren Knoten unter eine zuvor festgelegte Mindestanzahl fällt (Aggarwal, 2015, S. 100; Tan et al., 2020, S. 157).

Abbildung 2 zeigt einen binären Entscheidungsbaum angewandt auf die in dieser Arbeit unter anderem herangezogenen volkswirtschaftlichen Daten „Handelsbilanz“ und „Einzelhandelsumsätze“. Als Wurzelknoten (grau hinterlegt) fungiert in diesem Beispiel die „monatl. Veränderung der Handelsbilanz“ deren Ausprägung wie bereits oben erwähnt das Entscheidungskriterium darstellt. Die Entscheidungen an den inneren Knoten (hellgrau hinterlegt) in den Ebenen 2 und 3 werden ebenfalls anhand der „monatl. Veränderung der Handelsbilanz“ sowie der „monatl. Veränderung der Einzelhandelsumsätze“ getroffen. In den Blattknoten (weiß hinterlegt) erfolgt schlussendlich die Zuordnung zu den Klassen „Aktienindex steigt“ und „Aktienindex fällt“.

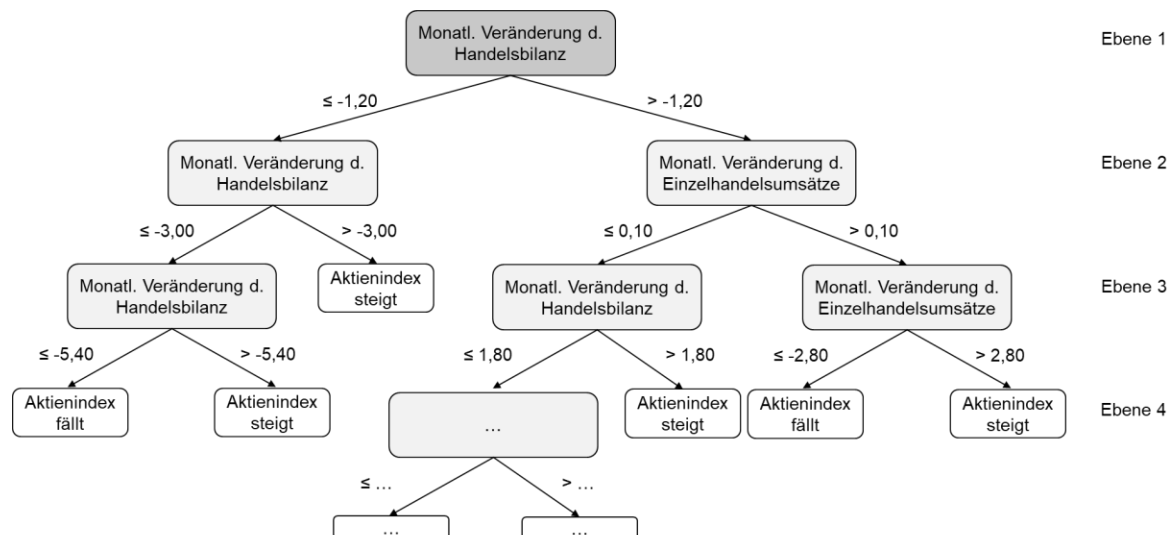


Abbildung 2: Binärer Entscheidungsbaum mit den Attributen Handelsbilanz und Einzelhandelsumsätze

Um die Beobachtungen im Wurzelknoten und jedem inneren Knoten aufzuspalten, muss ein geeignetes Attribut sowie ein geeigneter Attribute-Wert gewählt werden. Dabei ist es

wichtig, den Typ der Attribute zu kennen (Shafer et al., 1996, S. 546). Die in dieser Arbeit verwendeten Daten sind kontinuierlicher Natur, beispielsweise der Manufacturing PMI, welcher jeden Wert zwischen 0 und 100 annehmen kann (IHS Markit, o.J.). Die Methodik, um das am besten geeignete Attribut sowie den am besten geeigneten Wert des Attributs für ein Aufspalten der Daten (Verzweigung im Entscheidungsbaum) zu finden, wird nachfolgend erläutert. Dazu kann eine Bedingung wie folgt formuliert werden:

$$A \leq w \text{ oder } A > w, \quad (2)$$

wobei A das ausgewählte Attribut beziehungsweise Attribut-Variable und w den Wert bezeichnet, an dem die Beobachtungen aufgespalten werden (Rastogi and Shim, 2000, S. 319f). In Abbildung 3 wird beispielsweise das Attribut „monatliche Veränderung der Einzelhandelsumsätze“ („ A “) herangezogen, wobei der am besten geeignete Wert zur Aufspaltung bei $-0,65$ („ w “) liegt. Dementsprechend würden alle Beobachtungen in diesem Fall wie folgt aufgespalten werden:

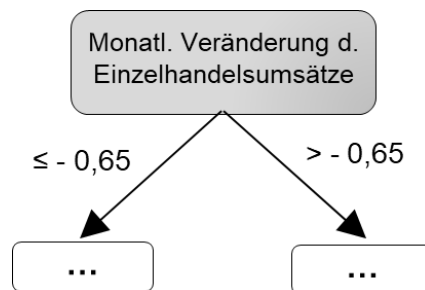


Abbildung 3: Binäre Entscheidungssituation

Um nun das am besten geeignete Attribut und den am besten geeigneten Wert für die Aufspaltung auszuwählen, lassen sich zahlreiche Möglichkeiten in der Literatur finden. Der Grundgedanke hinter diesen Methoden ist, disjunkte Mengen zu erhalten, in denen möglichst nur Beobachtungen einer Klasse zu finden sind. Das bedeutet, der Attribute-Test sollte die Beobachtungen möglichst so aufspalten, dass auf jedem neu entstehenden Ast des Entscheidungsbaumes nur Beobachtungen derselben Klasse zu finden sind (Shafer et al., 1996, S. 545-547). Eine Spaltungs-Bedingung wird folglich dann als wünschenswert erachtet, wenn daraus disjunkte Teilmengen resultieren, in welchen die Beobachtungen jeweils möglichst einer Klasse entsprechen. Entsprechen alle

Beobachtungen einer Klasse, ist folglich keine weitere Aufspaltung mehr notwendig beziehungsweise möglich (Tan et al., 2020, S. 147).

In nachfolgender Abbildung 4 ist das Ergebnis einer beispielhaften binären Aufspaltung für zwei verschiedene Werte des Attributs „monatliche Veränderung der Einzelhandelsumsätze“ dargestellt. Es ist ersichtlich, dass bei einem Wert von -0,65 ein Aufspalten dazu führt, dass im linken Ast alle Beobachtungen der positiven Klasse entsprechen. Beim Aufspalten anhand eines Werts von -0,50 entsprechen im linken Ast sieben Beobachtungen der positiven Klasse und eine Beobachtung der negativen Klasse. Betrachtet man nun nur die linken Äste, gelingt es mithilfe des Werts von -0,65 besser die Beobachtungen nach der Klasse zu spalten.

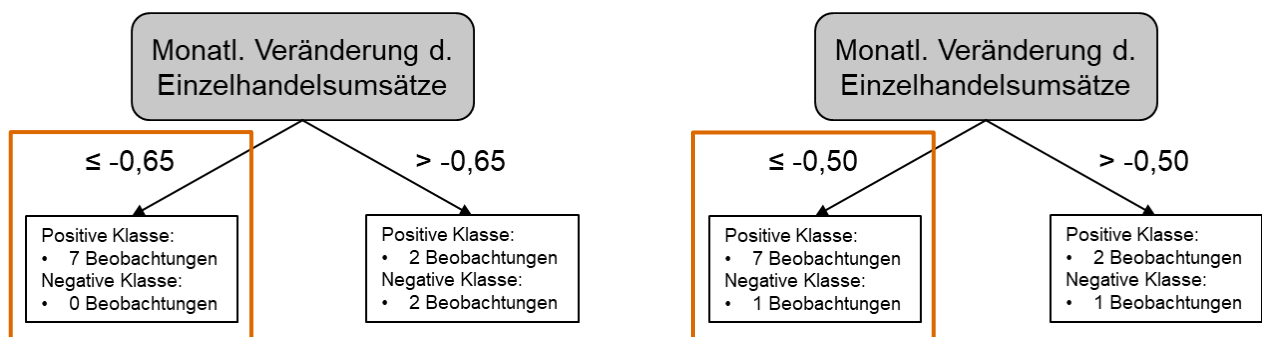


Abbildung 4: Aufspaltung anhand unterschiedlicher Attribute-Werte

Für die Auswahl der am besten geeigneten Aufspalte-Bedingung werden häufig sogenannte Maße der Unreinheit der Daten verwendet. Ein beliebtes Maß der Unreinheit der Daten stellt dabei der „Gini“-Index dar:

$$Gini = 1 - \sum_{i=0}^{c-1} p_i(k)^2, \quad (3)$$

wobei $p_i(k)$ dem Anteil der Beobachtungen der Klasse i an jedem Knoten k und c der Anzahl der Klassen entspricht. Jene Attribute-Test-Bedingung, die den geringsten gewichteten Gini-Index aufweist, wird für die Aufspaltung der Beobachtungen im getesteten Knoten herangezogen (Shafer et al., 1996, S. 547; Tan et al., 2020, S. 147f, 152). Nachfolgende Tabelle 1 stellt den Berechnungsvorgang am Beispiel der monatlichen Veränderung der Einzelhandelsumsätze in Deutschland dar. Dabei ist ersichtlich, dass der Entscheidungsbaum in diesem Knoten bei einem Wert von -0,65

aufgespalten wird, da dieser Wert den geringsten gewichteten Gini-Index in Höhe von 0,18 aufweist.

(1) Klasse	(2) Beob. Wert	(3) Spaltungsposition	(4) # positive Klasse		(5) # negative Klasse		(6) Gini-Index Teil 1	(7) Gini-Index Teil 2	(8) Gewichteter Gini-Index
			≤	>	≤	>			
Aktienindex steigt	-2,90	-2,95	0	9	0	2	N/A	0,30	N/A
Aktienindex steigt	-2,80	-2,85	1	8	0	2	0,00	0,32	0,29
Aktienindex steigt	-1,60	-2,20	2	7	0	2	0,00	0,35	0,28
Aktienindex steigt	-1,40	-1,50	3	6	0	2	0,00	0,38	0,27
Aktienindex steigt	-0,90	-1,15	4	5	0	2	0,00	0,41	0,26
Aktienindex steigt	-0,90	-0,90	6	3	0	2	0,00	0,48	0,22
Aktienindex steigt	-0,70	-0,80	6	3	0	2	0,00	0,48	0,22
Aktienindex fällt	-0,60	-0,65	7	2	0	2	0,00	0,50	0,18
Aktienindex fällt	-0,40	-0,50	7	2	1	1	0,22	0,44	0,28
Aktienindex steigt	-0,30	-0,35	7	2	2	0	0,35	0,00	0,28
Aktienindex steigt	-0,30	-0,30	9	0	2	0	0,30	N/A	N/A

Tabelle 1: Berechnungsbeispiel Gini-Index

Die „Klasse“ in Tabelle 1 gibt an, ob die prozentuale Veränderung des Aktienindex bei dieser Beobachtung positiv oder negativ ist, dementsprechend gibt sie an, ob der Kurs des Aktienindex steigt oder fällt.

Die „beobachteten Werte“ sind die aufsteigend geordneten Ausprägungen der monatlichen Veränderung der Einzelhandelsumsätze. Die „Spaltungsposition“ ist der Durchschnitt aus zwei aufeinanderfolgenden, beobachteten Werten (Shafer et al., 1996, S. 547; Tan et al., 2020, S. 153), wobei für den ersten Wert (- 2,90) ein Wert kleiner als - 2,90 zusätzlich zur Durchschnittsbildung herangezogen wird. Diese Durchschnittsbildung wird gewählt, da sonst willkürlich einer der zwei Werte, die in die Durchschnittsbildung einfließen, herangezogen werden müsste. Die Spalten „# positive Klasse“ und „# negative Klasse“ geben an, wie viele Beobachtungen der jeweiligen Klasse kleiner oder gleich dem Wert der Spaltungsposition sind, beziehungsweise wie viele Beobachtungen größer als dieser Wert sind. Dementsprechend sind in der hellgrün markierten Zeile sieben Beobachtungen der Veränderung der Einzelhandelsumsätze kleiner oder gleich - 0,65 Werteinheiten und gleichzeitig der positiven Klasse

(„Aktienindex steigt“) zuzuordnen. Der Gini-Index wird anschließend für jede der beiden Ungleichungen („ $\leq$ “ & „ $>$ “) nach (3) ermittelt, wobei sich für die „ $\leq$ “-Ungleichung ein Gini-Index von:

$$1 - \left(\frac{7}{7}\right)^2 - \left(\frac{0}{7}\right)^2 = 0, \quad (4)$$

für die „ $>$ “-Ungleichung ein Gini-Index von:

$$1 - \left(\frac{2}{4}\right)^2 - \left(\frac{2}{4}\right)^2 = 0,5 \quad (5)$$

sowie ein gewichteter Gini-Index in Anlehnung an Tan et al. (2020, S. 152) von:

$$0 * \frac{(7 + 0)}{(7 + 2 + 0 + 2)} + 0,5 * \frac{(2 + 2)}{(7 + 2 + 0 + 2)} = 0,18 \quad (6)$$

ergibt.

Ist nach der Aufspaltung der Beobachtungen in diesem Knoten keine der oben angeführten Stopp-Bedingungen erfüllt, werden die entstandenen Teilmengen analog erneut aufgespalten (Tan et al., 2020, S. 142f). Aus Tabelle 1 kann geschlossen werden, dass alle Beobachtungen, welche die Bedingung „ $\leq -0,65$ “ erfüllen, der positiven Klasse zugeordnet werden können. Folglich endet der Prozess hier in einem Blattknoten, welchem die positive Klasse zugeordnet wird. Jene Beobachtungen, die die Bedingung „ $> -0,65$ “ erfüllen, können jeweils zur Hälfte der positiven sowie der negativen Klasse („Aktienindex fällt“) zugeordnet werden.

Ist die gewählte Mindestanzahl an Beobachtungen nicht erreicht, führt der Algorithmus nun einen neuerlichen Test durch, um die verbleibenden Beobachtungen erneut aufzuspalten (Aggarwal, 2015, S. 100; Tan et al., 2020, S. 157).

3.2.2 Kritische Würdigung des Entscheidungsbaum-Klassifikators

Wie in Unterabschnitt 3.2.1 dargelegt, gibt es mehrere Stellschrauben, an denen bei der Konstruktion eines Entscheidungsbaumes gedreht werden kann.

Eine wichtige Entscheidung betrifft dabei die Auswahl eines geeigneten Maßes für die Unreinheit der Daten – in dieser Arbeit wurde beispielsweise der Gini-Index herangezogen. Es ist aber durchaus möglich auch andere Maße wie etwa den „Information Gain“ (Karaca and Cattani, 2019, S. 178f) oder eine Weiterentwicklung dieses Maßes, das „Gain Ratio“ heranzuziehen (Karaca and Cattani, 2019, S. 195).

Wie erwähnt, ist es auch notwendig, festzulegen, wann der Entscheidungsbaumwachstumsprozess gestoppt werden soll. Dies kann unter anderem erfolgen, wenn alle Beobachtungen in einem Knoten dieselbe Klasse aufweisen, die Beobachtungen alle dieselben Attribute-Werte aufweisen oder aber wenn die Anzahl an Beobachtungen nach einem weiteren Spaltungs-Vorgang eine festgelegte Mindestanzahl unterschreitet (Aggarwal, 2015, S. 100; Tan et al., 2020, S. 142f, 157).

Wie bereits auf den ersten Blick ersichtlich, ist die Wahl des geeigneten Maßes für die Unreinheit der Daten sowie die Entscheidung, wann der Wachstumsprozess des Entscheidungsbaumes beendet werden soll, mit sehr vielen subjektiven Entscheidungen verbunden. Des Weiteren kann festgestellt werden, dass die Anwendung von Entscheidungsbäumen mit steigender Attribute-Anzahl sehr rechenintensiv werden kann, da in jedem Knoten alle möglichen Varianten – das bedeutet, alle Attribute und potentiellen Aufspaltungspositionen der Attribute – zum Aufspalten der Beobachtungen getestet werden müssen (Aggarwal, 2015, S. 114).

Um wiederum der Intention der vorliegenden Arbeit, eine einfache, mit objektiven Entscheidungen verbundene, ohne viel rechnerischem Aufwand umsetzbare Anlagestrategie zu finden, Rechnung zu tragen, wird der Entscheidungsbaum als Generator für Handelssignale aufgrund der eben genannten Kritikpunkte in den nachfolgenden Analysen nicht weiter betrachtet.

3.2.3 Regel-basierter Klassifikator

Eine weitere Möglichkeit zur Klassifizierung von Daten stellen regel-basierte Klassifikatoren dar. Diese verwenden zur Einteilung von Beobachtungen in verschiedene Klassen eine Vielzahl von „wenn-dann“-Bedingungen bzw. Regeln, die zusammen eine Art Regelwerk bilden. Jede Regel dieses Regelwerks kann wie folgt dargestellt werden:

$$R: \text{Bedingung} \rightarrow \text{Klasse}, \quad (7)$$

wobei in die *Bedingung* die Inputs für den Klassifikator fließen. Darin können ein oder mehrere Attribute eingebunden werden, die jeweils für sich wiederum bestimmte Kriterien erfüllen müssen (Han et al., 2012, S. 355f). Beispielsweise kann eine Bedingung wie folgt formuliert werden:

$$\begin{aligned} \text{Handelsbilanz (1.Differenz)} \leq -1,20 \wedge \text{Einzelhandelsumsätze} \\ (\text{monatl.Veränderung}) \leq 0,60 \end{aligned} \quad (8)$$

Dabei muss jedes Attribut in der Bedingung eine festgelegte Voraussetzung erfüllen. Die einzelnen Attribute können mit den jeweiligen, für sie geltenden, kritischen Werten über logische Operatoren verbunden werden. Wie in (8) dargestellt, muss beispielsweise die erste Differenz der Handelsbilanz kleiner oder gleich -1,20 Werteinheiten und die monatliche Veränderung der Einzelhandelsumsätze kleiner oder gleich 0,60 Werteinheiten sein, damit die Bedingung und somit auch die gesamte Regel überhaupt zum Einsatz kommen. Die *Klasse* auf der rechten Seite in (7) wird als Ergebnis der Regel ausgegeben (Tan et al., 2020, S. 397f). Regel R_b ergibt sich, wie in (9) dargestellt, aus der Bedingung (8). Den gegebenen Attribute-Ausprägungen wird die Klasse „Aktienindex fällt“ zugeordnet:

$$\begin{aligned} R_b: \text{Handelsbilanz (1.Differenz)} \leq -1,20 \wedge \text{Einzelhandelsumsätze} \\ (\text{monatl.Veränderung}) \leq 0,60 \rightarrow \text{Aktienindex fällt} \end{aligned} \quad (9)$$

Damit ein Regelwerk funktionieren kann, müssen für die darin enthaltenen Regeln einige Kriterien gelten. Jede Beobachtung sollte nur maximal eine Regel im Regelwerk erfüllen. Regeln, die diese Eigenschaft besitzen, werden auch als sich gegenseitig ausschließende Regeln bezeichnet (Tan et al., 2020, S. 400). Für den Fall, dass

Beobachtungen mehrere Regeln eines Regelwerks erfüllen, bietet sich als Möglichkeit an, die Regeln im Regelwerk entsprechend zu ordnen und Beobachtungen anhand einer festgelegten Rangordnung zu klassifizieren. Diese Vorgehensweise soll verhindern, dass zwei Regeln, die beispielsweise verschiedene Klassen zuordnen, zugleich ausgewählt beziehungsweise erfüllt werden (Han et al., 2012, S. 356f). Die Ordnung kann, wie in der vorliegenden Arbeit dargestellt, beispielsweise anhand der vorhergesagten Klasse durchgeführt werden. Regeln für die einzelnen Klassen werden hierbei im Regelwerk gruppiert dargestellt, wobei es irrelevant ist, welche konkrete Regel angesprochen wird, solange eine Regel der entsprechenden Klasse angesprochen wird (Aggarwal, 2015, S. 125f).

Des Weiteren sollte ein Regelwerk umfassend sein. Das bedeutet, es sollte für jede Bedingung und entsprechend für jede Kombination von Attribute-Tests analog zu (8) zumindest eine Regel im Regelwerk erfüllt sein (Tan et al., 2020, S. 400). Wenn das nicht der Fall ist, kann mit einer Standardregel Abhilfe geschaffen werden. Dabei ordnet diese Standardregel allen Beobachtungen, die keine Regel erfüllen, eine im Vorhinein festgelegte Klasse zu. Diese Klasse entspricht dabei zumeist der am häufigsten vorliegenden Klasse aller Beobachtungen (die im Lernprozess berücksichtigt werden), die keine Regel erfüllen (Han et al., 2012, S. 357).

Um die entsprechenden Regeln, aus denen das Regelwerk besteht, aufzustellen, müssen diese aus den vorhandenen Daten generiert werden. Dabei unterscheidet man zwischen Methoden der direkten und Methoden der indirekten Regelextraktion, wobei für die vorliegende Arbeit die Methodik der direkten Regelextraktion gewählt wird. Für die Methode der indirekten Extraktion der Regeln müssten Regeln aus anderen Klassifizierungsmodellen extrahiert werden (Tan et al., 2020, S. 401), was wiederum entgegen der Zielsetzung der vorliegenden Arbeit wäre. Dementsprechend wird in weiterer Folge auf die direkte Methode eingegangen.

Wie bereits oben erwähnt, werden die Regeln nach der vorhergesagten Klasse geordnet und folglich auch nach der vorhergesagten Klasse aus den vorhandenen Daten extrahiert. Um zu entscheiden für welche Klasse die Regeln zuerst extrahiert werden,

wird die Häufigkeit der jeweiligen Klasse („Aktienindex steigt“, „Aktienindex fällt“) herangezogen (Tan et al., 2020, S. 401). So ist es beispielsweise möglich, Regeln für die seltenere Klasse zuerst zu extrahieren, um sicherzustellen, dass auch für diese Klasse Regeln abgeleitet werden (Aggarwal, 2015, S. 125; Tan et al., 2020, S. 401). In weiterer Folge wird davon ausgegangen, dass die Klasse „Aktienindex steigt“ in den zur Verfügung stehenden Beobachtungen vorherrschend ist. Dementsprechend wird mit der Regelerstellung für die Klasse „Aktienindex fällt“ begonnen.

Da es nur zwei Klassen gibt, müssen für die Klasse „Aktienindex steigt“ keine separaten Regeln extrahiert werden (Tan et al., 2020, S. 401).

In einem ersten Schritt ist es das Ziel, jene Regel aus den vorhandenen Daten zu extrahieren, welche die Klasse „Aktienindex fällt“ am besten beschreibt. Dabei sollten möglichst viele der Beobachtungen dieser Klasse und möglichst wenige der Klasse „Aktienindex steigt“ entsprechen. Eine Regel wird dabei auf Basis der besten Attribute-Bedingung erstellt und gegebenenfalls um weitere Attribute-Bedingungen ergänzt, solange bis keine Verbesserung mehr erzielt werden kann (Han et al., 2012, S. 359-361). Mögliche Kriterien, um festzustellen, welche Attribute-Bedingung zu Beginn in eine Regel inkludiert werden sollte, beziehungsweise um festzulegen, ob eine weitere Attribute-Bedingung in eine bereits vorhandene Regel eingebettet werden sollte, stellen beispielsweise die Metriken „Genauigkeit“ sowie die „Laplace-Metrik“, eine Erweiterung, welche auch die Abdeckung der Regel in Bezug auf die Datenbasis miteinbezieht, dar. Dabei weisen höhere Werte generell auf eine Verbesserung hin (Tan et al., 2006, S. 216-218).

Die Genauigkeitsmetrik wird wie folgt ermittelt:

$$\text{Genauigkeitsmetrik} = \frac{n_f}{n}, \quad (10)$$

wobei n_f die Anzahl der Beobachtungen mit der Klasse „Aktienindex fällt“, welche die Regel erfüllen und n die Anzahl der Beobachtungen, sowohl mit der Klasse „Aktienindex steigt“ als auch mit der Klasse „Aktienindex fällt“, die die Regel erfüllen, angeben (Han et al., 2012, S. 356). Die Laplace-Metrik wird wie folgt ermittelt:

$$\text{Laplace} - \text{Metrik} = \frac{n_f + 1}{n + c}, \quad (11)$$

wobei n und n_f analog zu Gleichung (10) zu sehen sind und c die Anzahl der möglichen Klassen – in der vorliegenden Arbeit zwei – angibt (Clark and Boswell, 1991, S. 153). Zusätzlich kann noch der Abdeckungsgrad der Regel bestimmt werden, welcher zwischen 0 % und 100 % liegen kann. Dieser wird wie folgt ermittelt:

$$\text{Abdeckungsgrad} = \frac{n}{n_g}, \quad (12)$$

wobei n analog zu Gleichung (10) zu sehen ist und n_g die Gesamtanzahl der Beobachtungen, unabhängig davon ob diese die Regel erfüllen, angibt (Han et al., 2012, S. 356).

Nachfolgend werden in Tabelle 2 und Tabelle 3 die dargestellten Schritte beispielhaft für die monatliche Veränderung der Einzelhandelsumsätze sowie für die erste Differenz der Handelsbilanz dargestellt:

Einzelhandelsumsätze (monatliche Veränderung)						Handelsbilanz (1. Differenz)					
(1) Wertgrenze	(2) # positive Klasse (≥ Wertgrenze)	(3) # negative Klasse (≥ Wertgrenze)	(4) Genauigkeits-Metrik	(5) Laplace-Metrik	(6) Abdeckungs-Grad	(1) Wertgrenze	(2) # positive Klasse (≥ Wertgrenze)	(3) # negative Klasse (≥ Wertgrenze)	(4) Genauigkeits-Metrik	(5) Laplace-Metrik	(6) Abdeckungs-Grad
-3,20	28	20	0,42	0,42	100%	-9,30	28	20	0,42	0,42	100%
-2,90	27	20	0,43	0,43	98%	-8,00	28	19	0,40	0,41	98%
-2,80	26	20	0,43	0,44	96%	-5,80	28	18	0,39	0,40	96%
-2,50	24	20	0,45	0,46	92%	-5,00	28	17	0,38	0,38	94%
-2,10	24	19	0,44	0,44	90%	-4,00	27	16	0,37	0,38	90%
-1,70	23	19	0,45	0,45	88%	-3,90	26	16	0,38	0,39	88%
-1,60	23	18	0,44	0,44	85%	-3,30	26	15	0,37	0,37	85%
-1,50	22	18	0,45	0,45	83%	-3,00	25	15	0,38	0,38	83%
-1,40	21	18	0,46	0,46	81%	-2,50	24	14	0,37	0,38	79%
...	...	...	...	...	...	...	...	...	...	...	...
0,80	6	8	0,57	0,56	29%	2,70	8	2	0,20	0,25	21%
1,00	5	7	0,58	0,57	25%	3,10	8	1	0,11	0,18	19%
1,20	4	7	0,64	0,62	23%	3,80	8	0	0,00	0,10	17%
1,30	3	7	0,70	0,67	21%	3,90	7	0	0,00	0,11	15%
1,40	3	5	0,63	0,60	17%	4,60	6	0	0,00	0,13	13%
1,50	3	4	0,57	0,56	15%	5,60	5	0	0,00	0,14	10%
1,90	2	3	0,60	0,57	10%	6,80	4	0	0,00	0,17	8%
2,50	1	3	0,75	0,67	8%	6,90	3	0	0,00	0,20	6%
3,10	1	1	0,50	0,50	4%	7,30	2	0	0,00	0,25	4%
6,30	0	1	1,00	0,67	2%	7,80	1	0	0,00	0,33	2%

Tabelle 2: Ermittlung der Metriken für $\geq$ Ungleichung

Einzelhandelsumsätze (monatliche Veränderung)						Handelsbilanz (1. Differenz)					
(1) Wertgrenze	(2) # positive Klasse ($\leq$ Wertgrenze)	(3) # negative Klasse ($\leq$ Wertgrenze)	(4) Genauigkeits-Metrik	(5) Laplace-Metrik	(6) Abdeckungs-Grad	(1) Wertgrenze	(2) # positive Klasse ($\leq$ Wertgrenze)	(3) # negative Klasse ($\leq$ Wertgrenze)	(4) Genauigkeits-Metrik	(5) Laplace-Metrik	(6) Abdeckungs-Grad
-3,20	1	0	0,00	0,33	2%	-9,30	0	1	1,00	0,67	2%
-2,90	2	0	0,00	0,25	4%	-8,00	0	2	1,00	0,75	4%
-2,80	4	0	0,00	0,17	8%	-5,80	0	3	1,00	0,80	6%
-2,50	4	1	0,20	0,29	10%	-5,00	1	4	0,80	0,71	10%
-2,10	5	1	0,17	0,25	13%	-4,00	2	4	0,67	0,63	13%
-1,70	5	2	0,29	0,33	15%	-3,90	2	5	0,71	0,67	15%
-1,60	6	2	0,25	0,30	17%	-3,30	3	5	0,63	0,60	17%
-1,50	7	2	0,22	0,27	19%	-3,00	4	6	0,60	0,58	21%
-1,40	10	2	0,17	0,21	25%	-2,50	4	7	0,64	0,62	23%
...	...	...	...	...	...	...	...	...	...	...	...
0,80	23	13	0,36	0,37	75%	2,70	20	19	0,49	0,49	81%
1,00	24	13	0,35	0,36	77%	3,10	20	20	0,50	0,50	83%
1,20	25	13	0,34	0,35	79%	3,80	21	20	0,49	0,49	85%
1,30	25	15	0,38	0,38	83%	3,90	22	20	0,48	0,48	88%
1,40	25	16	0,39	0,40	85%	4,60	23	20	0,47	0,47	90%
1,50	26	17	0,40	0,40	90%	5,60	24	20	0,45	0,46	92%
1,90	27	17	0,39	0,39	92%	6,80	25	20	0,44	0,45	94%
2,50	27	19	0,41	0,417	96%	6,90	26	20	0,43	0,44	96%
3,10	28	19	0,40	0,41	98%	7,30	27	20	0,43	0,43	98%
6,30	28	20	0,42	0,420	100%	7,80	28	20	0,42	0,42	100%

Tabelle 3: Ermittlung der Metriken für $\leq$ Ungleichung

In der Spalte „Wertgrenze“ in Tabelle 2 und Tabelle 3 werden die unterschiedlichen Ausprägungen der Beobachtungen angegeben. Dabei wird jede Ausprägung einmal als Wertgrenze herangezogen, öfters auftretende Ausprägungen fließen somit nur einmal als Wertgrenze in die Berechnung ein. Im Fall der Einzelhandelsumsätze (monatliche Veränderung) beträgt der minimale beobachtete Wert -3,20 und der maximale beobachtete Wert 6,30. Bei der Handelsbilanz (erste Differenz) betragen diese Werte - 9,30 beziehungsweise 7,80. Im Zuge der Berechnungen wurden 33 unterschiedliche Wertgrenzen für die Einzelhandelsumsätze (monatliche Veränderung) sowie 43 für die Handelsbilanz (erste Differenz) ermittelt und getestet. Zur Veranschaulichung werden in Tabelle 2 und Tabelle 3 allerdings nur je 19 Wertgrenzen dargestellt. An diesen Wertgrenzen wird in weiterer Folge in den Spalten „# positive Klasse ($\geq$ Wertgrenze)“ sowie „# negative Klasse ($\geq$ Wertgrenze)“ (analoge Vorgehensweise mit „ $\leq$ “ in Tabelle 3) die Häufigkeit der Beobachtungen mit positiver sowie negativer Klasse festgestellt. In Tabelle 2 sind folglich bei der monatlichen Veränderung der Einzelhandelsumsätze 28 Beobachtungen der positiven Klasse („Aktienindex steigt“) größer oder gleich dem Wert -3,20 und 20 Beobachtungen der negativen Klasse („Aktienindex fällt“) größer oder gleich dem Wert - 3,20. Daraus lassen sich nun anhand der Formeln (10), (11) und (12) die weiteren Kennzahlen berechnen.

Schlussendlich wird in der Entscheidungsfindung, welche Attribute-Bedingung zu Beginn in eine Regel inkludiert werden sollte beziehungsweise, um festzulegen, ob eine weitere

Attribute-Bedingung in eine bereits vorherrschende Regel eingebettet werden sollte, auf die Laplace-Metrik abgestellt, da diese, wie bereits oben erwähnt, verglichen mit der Genauigkeitsmetrik zusätzlich die Abdeckung der Regel in Bezug auf die Datenbasis berücksichtigt (Tan et al., 2006, S. 216-218).

Diese Tatsache ist beispielsweise gut in Tabelle 3 bei der ersten Differenz der Handelsbilanz erkennbar. Betrachtet man die beiden Wertgrenzen -8,00 und -5,80 lässt sich beobachten, dass bei diesen beiden Wertgrenzen eine Genauigkeitsmetrik in Höhe von 1,00 errechnet wird. Betrachtet man nun die Laplace-Metrik, lässt sich erkennen, dass diese im Fall der Wertgrenze -5,80 höher ist als bei der Wertgrenze -8,00 (0,80 beziehungsweise 0,75). Die Laplace-Metrik und die Genauigkeitsmetrik werden generell sehr ähnlich ermittelt, allerdings könnte man die Laplace-Metrik als eine um die Regelabdeckung bereinigte Genauigkeitsmetrik verstehen. Die Addition der Zahl 1 im Zähler sowie die Addition der Anzahl der Klassen c im Nenner, hat bei großer Abdeckung der Regel eine eher geringe Auswirkung auf den Wert und die Laplace-Metrik wird dementsprechend nahe an der Genauigkeitsmetrik liegen. Bei geringer Regelabdeckung werden diese Additionen allerdings eine große Auswirkung haben und der Wert der Laplace-Metrik wird sich stärker von der Genauigkeitsmetrik unterscheiden (Clark and Boswell, 1991, S. 152f).

Vergleicht man nun die Laplace-Metrik für alle durchgeführten Tests in Tabelle 2 und Tabelle 3, ergibt sich der höchste Wert mit rund 0,80 bei der ersten Differenz der Handelsbilanz und einem Wert von -5,80. Somit lautet der erste Teil der Bedingung in Regel 1:

$$\text{Handelsbilanz (1. Differenz)} \leq -5,80 \quad (13)$$

In einem nächsten Schritt wird nun versucht, die Regel durch Erweiterungen in der Test-Bedingung zu verbessern (Han et al., 2012, S. 360f). Dazu werden nur jene Beobachtungen herangezogen, die vom ersten Teil der Bedingung (siehe (13)) angesprochen werden – folglich verbleiben drei Beobachtungen, wie in Tabelle 4 dargestellt ist.

Einzelhandelsumsätze (monatliche Veränderung)						Einzelhandelsumsätze (monatliche Veränderung)					
(1) Wertgrenze	(2) # positive Klasse ($\geq$ Wertgrenze)	(3) # negative Klasse ($\geq$ Wertgrenze)	(4) Genauigkeits-Metrik	(5) Laplace-Metrik	(6) Abdeckungs-Grad	(1) Wertgrenze	(2) # positive Klasse ($\leq$ Wertgrenze)	(3) # negative Klasse ($\leq$ Wertgrenze)	(4) Genauigkeits-Metrik	(5) Laplace-Metrik	(6) Abdeckungs-Grad
0,60	0	3	1,00	0,80	100%	0,60	0	1	1,00	0,67	33%
0,70	0	2	1,00	0,75	67%	0,70	0	2	1,00	0,75	67%
2,50	0	1	1,00	0,67	33%	2,50	0	3	1,00	0,80	100%

Tabelle 4: Ermittlung der Metriken für verbleibende Beobachtungen

Wie in Tabelle 4 ersichtlich, kann die Laplace-Metrik durch die Ergänzung eines weiteren Bedingungsteils nicht erhöht werden. Somit besteht Regel 1 nun aus Ungleichung (13) und kann wie folgt in das Regelwerk übernommen werden:

$$R_1: \text{Handelsbilanz (1. Differenz)} \leq -5,80 \rightarrow \text{Aktienindex fällt} \quad (14)$$

Jene Beobachtungen, die Regel 1 ansprechen, werden im nächsten Schritt eliminiert, da sonst in jedem Durchlauf dieselbe Regel generiert werden würde, und für die verbleibenden Beobachtungen wird die beschriebene Prozedur erneut durchgeführt. Neue Regeln werden so lange generiert bis ein Anhalte-Kriterium erfüllt ist (Tan et al., 2020, S. 401f). Dieses kann unter anderem dann erfüllt sein, wenn keine Beobachtungen mehr für die betrachtete Klasse vorliegen (Aggarwal, 2015, S. 126). Für alle Beobachtungen, die am Ende dieses Prozesses keine Regel erfüllen, wird, wie bereits oben erwähnt, eine Standardregel eingeführt und die Klasse entsprechend der vorherrschenden Klasse aller verbleibenden Beobachtungen zugeordnet (Han et al., 2012, S. 357).

3.2.4 Kritische Würdigung des regel-basierten Klassifikators

Wie in Unterabschnitt 3.2.3 dargelegt, gibt es mehrere Stellschrauben, an denen bei der Konstruktion eines regel-basierten Klassifikators gedreht werden kann.

So ist es beispielsweise von zentraler Bedeutung, welche Metrik(en) herangezogen wird (werden), um zu entscheiden, welche Attribute-Bedingung zu Beginn in eine Regel inkludiert beziehungsweise ob eine weitere Attribute-Bedingung zu einer Regel hinzugefügt werden sollte. Dafür können neben den dargestellten Metriken auch andere wie beispielsweise der FOIL Information Gain oder die Likelihood-Ratio herangezogen werden (Han et al., 2012, S. 362f). Diese Entscheidung kann sehr subjektiv geprägt sein.

Analog zum Entscheidungsbaum-Klassifikator wird auch die Anwendung des regel-basierten Klassifikators mit steigender Attribute-Anzahl sehr rechenintensiv, da für jedes

Attribut die Tests, wie in Tabelle 2 und Tabelle 3 dargestellt, durchgeführt werden müssen. Schlussendlich kann auch die Wahl des Anhalte-Kriteriums von subjektiven Entscheidungen geprägt sein.

Um wiederum der Intention der vorliegenden Arbeit, eine einfache, mit objektiven Entscheidungen verbundene, ohne viel rechnerischem Aufwand umsetzbare Anlagestrategie zu finden, Rechnung zu tragen, wird der regel-basierte Klassifikator als Generator für Handelssignale aufgrund der eben genannten Kritikpunkte in den nachfolgenden Analysen nicht weiter betrachtet.

3.2.5 Klassifikator der k-nächsten Nachbarn („kNN“)

Eine weitere Möglichkeit, um Beobachtungen zu klassifizieren, stellt der Klassifikator der k-nächsten Nachbarn dar. Bei dieser Methode wird die Klasse einer Beobachtung im Wesentlichen aufgrund der Klasse deren nächster Nachbarn festgelegt (Aggarwal, 2015, S. 160). Diese Methode kann gut mit folgendem Spruch beschrieben werden: „If it walks like a duck, quacks like a duck, and looks like a duck, then it’s probably a duck.“ (Tan et al., 2020, S. 410). Im Unterschied zu den vorherigen beiden Methoden ist bei dieser Art der Klassifikation keine vorangehende Modellerstellung notwendig. Ein Modell wird erst dann erstellt, sobald eine zu klassifizierende Beobachtung vorliegt (Han et al., 2012, S. 423).

Dabei werden alle bekannten Beobachtungen in einem Raum mit n Dimensionen dargestellt, wobei n der Anzahl der Attribute entspricht. Sobald ein zu klassifizierender Datenpunkt vorliegt, wird über ein Abstandsmaß die Nähe des Datenpunktes zu den anderen Datenpunkten im n -dimensionalen Raum gemessen. Anschließend legt der Klassifikator die Klasse auf Basis der Mehrheit der k-nächsten Nachbarn des zu klassifizierenden Datenpunktes fest. Wird beispielsweise $k=3$ gewählt, so betrachtet der Algorithmus die drei Nachbarn, für die die geringsten Abstände gemessen werden. Auf Basis der Klasse dieser drei Nachbarn wird die Entscheidung zugunsten der Mehrheit getroffen (Karaca and Cattani, 2019, S. 273f). Weisen beispielsweise zwei Nachbarn die Klasse „Aktienindex steigt“ und ein Nachbar die Klasse „Aktienindex fällt“ auf, so wird dem neuen Datenpunkt die Klasse „Aktienindex steigt“ zugewiesen.

Als Abstandsmaß wird in weiterer Folge der Euklidische Abstand verwendet, welcher für n Dimensionen in Anlehnung an O'Brien (2014, S. 156) wie folgt ermittelt wird:

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2}, \quad (15)$$

wobei die x_i die Ausprägungen der Attribute eines zu klassifizierenden Datenpunktes und die y_i die jeweiligen Ausprägungen eines bekannten Datenpunktes darstellen.

Diese Vorgehensweise lässt sich nur auf numerische Daten anwenden und kann in dieser Arbeit herangezogen werden, da nur numerische Attribute Berücksichtigung finden (Karaca and Cattani, 2019, S. 274).

Vor Anwendung des Abstandsmaß (15) auf die verfügbaren Beobachtungen ist es wichtig, diese zu skalieren. Hat man beispielsweise zwei Attribute von denen eines Werte zwischen 0 und 100 und eines, welches Werte zwischen 0 und 1 annehmen kann, wird man feststellen, dass das Abstandmaß sehr durch das erste Attribut beeinflusst wird. Wird eine Skalierung nicht durchgeführt, ist das Abstandsmaß unter Umständen verzerrt (Tan et al., 2020, S. 414).

In der vorliegenden Arbeit wird auf die Standardisierung anhand des z-Werts als eine Form der Skalierung in Anlehnung an Wooldridge (2016, S. 657f) zurückgegriffen:

$$z_i = \frac{X_i - \mu}{\sigma}, \quad (16)$$

wobei, X_i den jeweiligen Attribute-Wert, μ den Mittelwert und σ die Standardabweichung der vorhandenen Beobachtungen darstellen. Die z_i sind schließlich durch einen Mittelwert von 0 und eine Standardabweichung von 1 charakterisiert.

Mithilfe der Standardisierung soll somit verhindert werden, dass ein Attribut, welches üblicherweise große Werte annimmt, die Berechnung des Abstandsmaßes dominiert (Tan et al., 2006, S. 64f).

Für den Mittelwert μ und die Standardabweichung σ werden in der vorliegenden Arbeit in der Folge das arithmetische Mittel sowie die Stichproben-Standardabweichung herangezogen.

In Tabelle 5 wird die beschriebene Methodik beispielhaft für die monatliche Veränderung der Einzelhandelsumsätze sowie die erste Differenz der Handelsbilanz angewandt.

(1) Beobachtung #	(2) Einzelhandelsumsätze (monatl. Veränderung)	(3) Standardisierte Einzelhandelsumsätze (monatl. Veränderung)	(4) Handelsbilanz (1. Differenz)	(5) Standardisierte Handelsbilanz (1. Differenz)	(6) Euklidischer Abstand (neue Beobachtung #11)	(7) Klasse
1	-0,30	-0,14	-1,00	-0,36	0,56	Aktienindex steigt
2	1,40	0,54	-1,80	-0,55	0,86	Aktienindex fällt
3	-0,30	-0,14	2,00	0,35	0,28	Aktienindex steigt
4	-2,10	-0,85	6,80	1,50	1,62	Aktienindex steigt
5	0,60	0,22	-8,00	-2,04	2,20	Aktienindex fällt
6	-2,80	-1,13	3,90	0,81	1,36	Aktienindex steigt
7	6,30	2,49	-2,10	-0,63	2,55	Aktienindex fällt
8	0,00	-0,02	-2,30	-0,67	0,84	Aktienindex fällt
9	-2,90	-1,17	1,40	0,21	1,24	Aktienindex steigt
10	0,40	0,14	5,60	1,21	1,05	Aktienindex steigt
Zu klassifizieren						
11	0,20	0,06	1,20	0,16	N/A	Aktienindex steigt
Mittelwert	0,05	0,00	0,52	0,00		
Standardabweichung	2,51	1,00	4,19	1,00		

Tabelle 5: Beispielhafte Darstellung des kNN-Klassifikators

Spalten 2 und 4 stellen die beobachteten Werte sowie deren Mittelwert und Standardabweichung dar. In den Spalten 3 und 5 werden die beobachteten Werte standardisiert. Dies erfolgt anhand der Gleichung (16). Dementsprechend ergibt sich der standardisierte Wert der Beobachtung #1 in Spalte 3 in Höhe von -0,14 aus:

$$\frac{(-0,30 - 0,05)}{2,51} \approx -0,14 \quad (17)$$

Der Euklidische Abstand zwischen der neuen, zu klassifizierenden Beobachtung #11 und Beobachtung #1 in Höhe von 0,56 ergibt sich analog zu Gleichung (15) aus:

$$\sqrt{(0,06 - (-0,14))^2 + (0,16 - (-0,36))^2} = 0,56 \quad (18)$$

In Tabelle 5 wird der Wert $k=3$ gewählt, folglich werden jene drei Beobachtungen mit den geringsten Abständen zur Beobachtung #11 als Referenz für die Festlegung der Klasse herangezogen. In diesem Fall sind das die Beobachtungen #1, #3 und #8 mit den dazugehörigen Klassen „Aktienindex steigt“, „Aktienindex steigt“ und „Aktienindex fällt“.

Da nach dem Mehrheitsprinzip entschieden wird, wird der Beobachtung #11 folglich die Klasse „Aktienindex steigt“ zugewiesen (Karaca and Cattani, 2019, S. 273f).

3.2.6 Kritische Würdigung des Klassifikators der k-nächsten Nachbarn

Auch beim Klassifikator der k-nächsten Nachbarn gibt es essentielle Parameter, die spezifiziert werden müssen, wobei subjektive Überlegungen in die Parameterwahl einfließen.

So ist es beispielsweise von zentraler Bedeutung in welcher Höhe k gewählt wird. Wird ein zu geringer Wert gewählt, kann der Klassifikator stark durch falsch klassifizierte Beobachtungen (misabeled data) beeinflusst werden. Wird k hingegen zu groß gewählt, so fließen weit von der zu klassifizierenden Beobachtung entfernte Datenpunkte in die Klassifikation mit ein und das Risiko einer Fehlklassifikation nimmt zu. Um dem Umstand der Beeinflussung durch zu weit entfernte Beobachtungen Rechnung zu tragen und den Einfluss von k zu reduzieren, bietet sich die Möglichkeit, eine abstandsgewichtete Entscheidung zu treffen, in welcher näher an der zu klassifizierenden Beobachtung liegende Datenpunkte einen größeren Einfluss bei der Klassifikationsentscheidung bekommen (Tan et al., 2020, S. 410-412). Je nachdem welche Methodik herangezogen wird, könnte im Ergebnis eine abweichende Klassifikationsentscheidung getroffen werden.

Eine weitere wichtige Entscheidung bei der Konstruktion des kNN-Klassifikators betrifft die Wahl des Abstandsmaßes wie Abu Alfeilat et al. (2019) zeigen. In der vorliegenden Arbeit wurde zur Veranschaulichung der Euklidische Abstand herangezogen. Aus Abu Alfeilat et al. (2019) lässt sich somit schließen, dass unterschiedliche Abstandsmaße ebenfalls zu unterschiedlichen Klassifikationsentscheidungen führen könnten, wobei eine Diskussion unterschiedlicher Abstandsmaße nicht Teil dieser Arbeit ist.

Um wiederum der Intention der vorliegenden Arbeit Rechnung zu tragen, wird der Klassifikator der k-nächsten Nachbarn als Generator für Handelssignale aufgrund der eben genannten, zu treffenden Konstruktionsentscheidungen in den nachfolgenden Analysen nicht weiter betrachtet.

3.2.7 Naive Bayes Klassifikator

Der letzte Klassifikator, welcher in der vorliegenden Arbeit betrachtet wird und welchem schlussendlich der Vorzug in den weiteren Analysen gegeben wird, ist der Naive Bayes Klassifikator.

Dieser gehört zur Gruppe der wahrscheinlichkeitsbasierten Klassifikatoren, die einer Beobachtung nicht nur die am besten passende Klasse zuordnen, sondern auch eine Wahrscheinlichkeit angeben, mit welcher die Beobachtung jeder verfügbaren Klasse zugeordnet werden kann. Dabei wird jene Klasse mit der höchsten Wahrscheinlichkeit häufig der Beobachtung zugeordnet (Aggarwal, 2015, S. 66).

Wie der Name des Klassifikators erahnen lässt, basiert dieser auf dem bekannten Satz von Bayes (Lewis, 1998, S. 4-6). Wie auch bei den vorangegangenen Analysen ist es hier das Ziel, einer Beobachtung aufgrund ihrer Attribute-Ausprägungen eine Klasse zuzuweisen. Der Satz von Bayes lautet wie folgt, wobei A die Klassenvariable und B die Attribute und deren Ausprägungen darstellen (Fahrmeir et al., 2016, S. 199):

$$P(A|B) = \frac{P(B|A) * P(A)}{P(B)} \quad (19)$$

Dabei werden die Wahrscheinlichkeit $P(A)$ als a-priori und $P(A|B)$ als a-posteriori Wahrscheinlichkeit bezeichnet. Diese Bezeichnungen ergeben sich aufgrund der Tatsache, dass $P(A)$ bereits vor Kenntnis der jeweiligen Ausprägungen in B bekannt ist und $P(A|B)$ erst nach Kenntnis von B ermittelt werden kann (Fahrmeir et al., 2016, S. 199). Beispielsweise kann $P(A=\text{Aktienindex steigt})$ als die Wahrscheinlichkeit eines ansteigenden beziehungsweise $P(A=\text{Aktienindex fällt})$ als die Wahrscheinlichkeit eines fallenden Aktienindex in allen Beobachtungen definiert werden.

Die a-posteriori Wahrscheinlichkeit oder der „Posterior“, auf der linken Seite in Gleichung (19), kann außerdem wie folgt dargestellt werden:

$$Posterior = \frac{Likelihood * Prior}{Evidence}, \quad (20)$$

wobei „Likelihood“ als Indikator verstanden werden kann, der angibt, welche Klasse, ceteris paribus, eher der wahren Klasse entspricht, gegeben die Attribute-Ausprägungen.

Der „Prior“ gibt, wie oben erläutert, die Wahrscheinlichkeit des Eintretens einer jeden Klasse in den Beobachtungen an. „Evidence“ ist eine Konstante, die sicherstellt, dass die a-posteriori Wahrscheinlichkeiten in Summe 1 ergeben (Duda et al., 2001, S. 20-22).

Basierend auf den vorhandenen Beobachtungen kann in der Trainingsphase mit den Attribute-Ausprägungen B in weiterer Folge die Likelihood $P(B|A)$, jene „Wahrscheinlichkeit“ des Eintretens von B gegeben die Klasse A , ermittelt werden. Anschließend kann mit dem Satz von Bayes die a-posteriori Wahrscheinlichkeit $P(A|B)$ für jede Klasse A ermittelt werden und jene Klasse, für welche diese am höchsten ist, der bisher unbekannten Beobachtung zugeordnet werden (Aggarwal, 2015, S. 66-69). Der Evidence-Faktor $P(B)$ kann bei der Ermittlung der a-posteriori Wahrscheinlichkeiten außer Acht gelassen werden, da $P(B)$ wie bereits oben erläutert eine Konstante ist, die sicherstellt, dass die a-posteriori Wahrscheinlichkeiten in Summe 1 ergeben und somit die Auswahl nicht beeinflusst (Duda et al., 2001, S. 22; Tan et al., 2020, S. 422).

Bei der Ermittlung von $P(B|A)$ kommt nun die naive Anwendung des Satz von Bayes zum Einsatz. Sei B beispielsweise ein Vektor bestehend aus einer Vielzahl von Attributen B_i , so müssten für die Ermittlung von $P(B|A)$ sehr viele unterschiedliche Parameter ermittelt werden, um für jede mögliche Kombination der Attribute mit deren zugehörigen Ausprägungen, gegeben eine Klasse A_i , eine (aussagekräftige) Likelihood zu erhalten. Außerdem würde diese Vorgehensweise eine sehr große Anzahl an Beobachtungen voraussetzen, die in den allermeisten Fällen nicht verfügbar ist. Aus diesem Grund wird beim Naive Bayes Klassifikator die naive Annahme getroffen, dass die Attribute untereinander unabhängig sind und nur von der Klasse A abhängen (Aggarwal, 2015, S. 670; Lewis, 1998, S. 5f). Der Klassifikator kann unter dieser Annahme wie folgt dargestellt werden:

$$P(A|B) = \frac{P(A) * \prod_{i=1}^d P(B_i|A)}{P(B)}, \quad (21)$$

wobei B aus d Attributen besteht. Anders als beim Satz von Bayes muss bei Naive Bayes nur die Likelihood jedes einzelnen Attributes (gegeben die Klassenvariable) ermittelt werden und nicht die jeder möglichen Kombination des Attribute-Sets B mit der Klassenvariablen (Aggarwal, 2015, S. 70; Tan et al., 2020, S. 421).

Um die Likelihood eines jeden Attributes $P(B_i = b_i | A = A_j)$ zu ermitteln, können bei kontinuierlichen Daten entweder Intervalle erstellt werden, oder es wird eine Wahrscheinlichkeitsverteilung zur Hilfe genommen (Tan et al., 2020, S. 423f). In der vorliegenden Arbeit wird die zweite Methode herangezogen, da bei dieser Methode die Wahl der „richtigen“ Grenzen für die Intervalle wegfällt.

Dabei wird in der Praxis häufig auf die Normalverteilung zurückgegriffen (Duda et al., 2001, S. 31; John and Langley, 1995, S. 339). Die Likelihoods beziehungsweise die Dichten werden nach dieser Methode mithilfe der Dichtefunktion der Verteilung ermittelt (Duda et al., 2001, S. 31). Die Dichtefunktion der Normalverteilung lautet allgemein wie folgt (Fahrmeir et al., 2016, S. 83):

$$f(x|\mu, \sigma) = \frac{1}{\sigma * \sqrt{2 * \pi}} * e^{-\frac{1}{2} * \left(\frac{x-\mu}{\sigma}\right)^2}, \quad (22)$$

wobei x den jeweiligen betrachteten Wert, μ den Mittelwert und σ die Standardabweichung darstellen. Für den Mittelwert μ wird in weiterer Folge das arithmetische Mittel $\bar{x}$ herangezogen. Für die Standardabweichung σ wird in weiterer Folge die Stichproben-Standardabweichung s herangezogen (Tan et al., 2020, S. 424).

In Anlehnung an Tan et al. (2020, S. 424) kann Formel (22) für die Ermittlung der Likelihoods beziehungsweise der Dichten in der vorliegenden Arbeit wie folgt formuliert werden:

$$P(B_i = b_i | A = A_j) = \frac{1}{\sigma_{ij} * \sqrt{2 * \pi}} * e^{-\frac{1}{2} * \left(\frac{b_i - \mu_{ij}}{\sigma_{ij}}\right)^2}, \quad (23)$$

wobei μ_{ij} dem arithmetischen Mittel und σ_{ij} der Stichproben-Standardabweichung der jeweiligen Klasse und des jeweiligen Attributes entsprechen. Die Parameter können dabei anhand der Beobachtungen des jeweiligen Attributes und Klasse ermittelt werden. Anhand der abgeleiteten Funktionen je Klasse kann in einem weiteren Schritt die Likelihood der Klassenzugehörigkeit einer neuen Beobachtung für jedes Attribut ermittelt und in Gleichung (21) angewandt werden.

Im Folgenden wird der Naive Bayes Klassifikator beispielhaft angewandt, wobei abermals die Attribute „monatliche Veränderung der Einzelhandelsumsätze“ sowie „erste Differenz

der Handelsbilanz“ herangezogen werden. In Tabelle 6 sind die Ausprägungen der zwei Attribute sowie des Klassenattributes dargestellt. Daraus lassen sich die a-priori Wahrscheinlichkeiten der beiden Klassen ermitteln. Diese betragen für die Klasse „Aktienindex steigt“:

$$P(A = \text{Aktienindex steigt}) = \frac{6}{10} = 0,60 \quad (24)$$

sowie für die Klasse „Aktienindex fällt“:

$$P(A = \text{Aktienindex fällt}) = \frac{4}{10} = 0,40 \quad (25)$$

Des Weiteren sind in Tabelle 6 der Mittelwert sowie die Stichproben-Standardabweichung für die beiden Attribute jeweils für die beiden Klassen dargestellt.

(1) Beobachtung #	(2) Einzelhandelsumsätze (monatl. Veränderung)	(3) Handelsbilanz(1. Differenz)	(4) Klasse
1	-0,30	-1,00	Aktienindex steigt
2	1,40	-1,80	Aktienindex fällt
3	-0,30	2,00	Aktienindex steigt
4	-2,10	6,80	Aktienindex steigt
5	0,60	-8,00	Aktienindex fällt
6	-2,80	3,90	Aktienindex steigt
7	6,30	-2,10	Aktienindex fällt
8	0,00	-2,30	Aktienindex fällt
9	-2,90	1,40	Aktienindex steigt
10	0,40	5,60	Aktienindex steigt
Mittelwert (Klasse „Aktienindex steigt“)	-1,33	3,12	
Mittelwert (Klasse „Aktienindex fällt“)	2,08	-3,55	
Standardabweichung (Klasse „Aktienindex steigt“)	1,44	2,88	
Standardabweichung (Klasse „Aktienindex fällt“)	2,87	2,97	

Tabelle 6: Ermittlung Mittelwert und Stichproben-Standardabweichung je Klasse

Betrachtet man beispielsweise das Attribut „monatliche Veränderung der Einzelhandelsumsätze“, ergibt sich der Mittelwert der Attribute-Ausprägungen, die der Klasse „Aktienindex steigt“ entsprechen – Beobachtungen #1, #3, #4, #6, #9, #10 – in Höhe von -1,33. Die Stichproben-Standardabweichung für dieselben Beobachtungen beträgt 1,44. Im nächsten Schritt werden mithilfe der ermittelten Parameter in Tabelle 6 und Gleichung (23) die Likelihoods für die Attribute (unter gegebenen Klassenvariablen), ermittelt. Die Ergebnisse sind in Tabelle 7 für die bisher unbekannten Beobachtungen #11 bis #15 dargestellt.

(1) Zu klassifizierende Datenpunkte	(2) Einzelhandels- umsätze (monatl. Veränderung)	(3) Handelsbilanz (1. Differenz)	(4) $P(EHU=b_{ij} A=$ Aktienindex steigt) („C“)	(5) $P(HB=b_{ij} A=$ Aktienindex steigt) („D“)	(6) $P(EHU=b_{ij} A=$ Aktienindex fällt) („E“)	(7) $P(HB=b_{ij} A=$ Aktienindex fällt) („F“)	(8) $C \cdot D \cdot 0,6$ („G“)	(9) $E \cdot F \cdot 0,4$ („H“)
11	1,00	-4,00	0,07434	0,00655	0,12941	0,13263	0,00029	0,00687
12	0,20	1,20	0,15712	0,11099	0,11219	0,03746	0,01046	0,00168
13	2,90	-3,20	0,00363	0,01251	0,13319	0,13323	0,00003	0,00710
14	-0,50	3,30	0,23459	0,13821	0,09292	0,00945	0,01945	0,00035
15	-2,30	3,80	0,22136	0,13465	0,04358	0,00633	0,01788	0,00011

Tabelle 7: Ermittlung der Likelihoods sowie Multiplikation von Likelihoods und Prior (Zählergröße)

In den Spalten 2 und 3 sind abermals die Merkmalsausprägungen der betrachteten Attribute dargestellt. In Spalten 4 bis 7 wird jeweils die Likelihood der Attribute, gegeben die Klasse, ermittelt. Dabei wird das Attribut „monatliche Veränderung der Einzelhandelsumsätze“ mit „EHU“ und das Attribut „erste Differenz der Handelsbilanz“ mit „HB“ bezeichnet. Beispielhaft ergeben sich für Beobachtung #11 in den Spalten 4 und 5 die Likelihoods bei gegebener Klasse „Aktienindex steigt“ nach Gleichung (23) in Höhe von rund 0,07 bzw. rund 0,01. Diese werden wie folgt ermittelt:

$$P(EHU = 1,00|A = \text{Aktienindex steigt}) = \frac{1}{1,44 * \sqrt{2 * \pi}} * e^{-\frac{1}{2} * \left(\frac{1,00 - (-1,33)}{1,44}\right)^2} \approx 0,07434 \quad (26)$$

$$P(HB = -4,00|A = \text{Aktienindex steigt}) = \frac{1}{2,88 * \sqrt{2 * \pi}} * e^{-\frac{1}{2} * \left(\frac{-4,00 - 3,12}{2,88}\right)^2} \approx 0,00655 \quad (27)$$

Die Werte in den Spalten 6 und 7 werden nach analoger Vorgehensweise für die Klasse „Aktienindex fällt“ ermittelt. In den Spalten 8 und 9 wird auf Basis der Gleichung (21) die Zählergröße (das bedeutet Likelihood * Prior) zur Berechnung der a-posteriori Wahrscheinlichkeit für beide Klassen ermittelt. Für Beobachtung #11 und Klasse „Aktienindex steigt“ ergibt sich somit analog zum Zähler (Likelihood * Prior) in Gleichung (21):

$$0,07434 * 0,00655 * 0,6 \approx 0,00029 \quad (28)$$

Für die Klasse „Aktienindex fällt“ ergibt sich die Zählergröße in Höhe von rund 0,00687 aus:

$$0,12941 * 0,13263 * 0,4 \approx 0,00687 \quad (29)$$

Anschließend werden diese, wie in Tabelle 8 ersichtlich, normiert, damit diese auf 100 Prozent aufsummiert werden können. Dies geschieht beispielsweise für die Klasse „Aktienindex steigt“ anhand der Formel:

$$\frac{0,00029}{0,00029 + 0,00687} \approx 0,04, \quad (30)$$

wobei der Nenner sicherstellt, dass die a-posteriori Wahrscheinlichkeiten in Summe 100 Prozent ergeben (Duda et al., 2001, S. 22).

(1) Zu klassifizierende Datenpunkte	(2) G	(3) H	(4) $P[DAX_{t+1} > DAX_t]$	(5) $P[DAX_{t+1} < DAX_t]$	(6) Klassifikations- entscheidung
11	0,00029	0,00687	4%	96%	Aktienindex fällt
12	0,01046	0,00168	86%	14%	Aktienindex steigt
13	0,00003	0,00710	0%	100%	Aktienindex fällt
14	0,01945	0,00035	98%	2%	Aktienindex steigt
15	0,01788	0,00011	99%	1%	Aktienindex steigt

Tabelle 8: Ermittlung der a-posteriori Wahrscheinlichkeiten

Zuletzt wird in Anlehnung an Aggarwal (2015, S. 66) jeder Beobachtung jene Klasse zugeordnet, die eine höhere normierte Wahrscheinlichkeit aufweist. Folglich wird den Beobachtungen #11 und #13 die Klasse „Aktienindex fällt“ und den Beobachtungen #12, #14 und #15 die Klasse „Aktienindex steigt“ zugewiesen.

3.2.8 Kritische Würdigung des Naive Bayes Klassifikators

Analog zu den vorangegangenen Klassifikatoren, wird in diesem Unterabschnitt die Anwendung des Naive Bayes Klassifikators kritisch gewürdigt.

Wie bereits oben erwähnt, kann bei kontinuierlichen Daten die Ermittlung der Likelihoods $P(B|A)$ über Intervalle oder über die Zuhilfenahme einer Wahrscheinlichkeitsverteilung erfolgen (Tan et al., 2020, S. 423f). In der vorliegenden Arbeit wird bewusst auf die zweite Vorgangsweise abgestellt, da auf diesem Wege einer (subjektiven) Entscheidung über die geeignete Auswahl der Intervalle vorgebeugt wird. Die Wahl der Normalverteilung als zugrundeliegende Verteilung könnte in gewisser Weise ebenfalls als subjektiv betrachtet werden, allerdings wird diese in der Praxis bei kontinuierlichen Daten häufig herangezogen (Duda et al., 2001, S. 31; John and Langley, 1995, S. 339). Dies kann in gewisser Weise durch den zentralen Grenzwertsatz, einem sehr elementaren Prinzip in der Statistik, untermauert werden. Der zentrale Grenzwertsatz sagt aus, dass sich die Verteilung von beispielsweise der Summe/Durchschnitte von Zufallsvariablen mit zunehmender Anzahl von Zufallsvariablen n , die in die Summenermittlung/

Durchschnittsermittlung einbezogen werden, der Normalverteilung annähert. Dabei ist es unerheblich welche Verteilung den Zufallsvariablen zugrunde liegt (Fahrmeir et al., 2016, S. 293-295; Wooldridge, 2016, S. 684).

Des Weiteren ist die Anwendung der Naive Bayes Methode bei einer großen Anzahl an Attributen geeignet (Aggarwal, 2015, S. 67). Anders als beispielsweise bei dem Entscheidungsbaum-Klassifikator oder bei dem regel-basierten Klassifikator ist eine steigende Attribute-Anzahl mit wenig zusätzlichem Rechenaufwand verbunden.

Ein weiterer Vorteil des Naive Bayes Klassifikators besteht darin, dass dieser nicht nur die jeweilige Klasse zuordnet sondern auch eine korrespondierende Wahrscheinlichkeit angibt, mit welcher die betrachtete Beobachtung einer Klasse zugeordnet werden kann (Aggarwal, 2015, S. 66).

3.3 Auswahl des anzuwendenden Klassifikators

Basierend auf den in Abschnitt 3.2 dargestellten Verfahren zur Klassifizierung und der kritischen Würdigung jeder Methodik wird in weiterer Folge der Naive Bayes Klassifikator für die Analysen der vorliegenden Arbeit herangezogen. Dieser wird unseres Erachtens am besten den Anforderungen an die Arbeit, eine einfache, mit möglichst objektiven Entscheidungen verbundene, ohne viel rechnerischem Aufwand umsetzbare Anlagestrategie zu finden, gerecht.

4 Überführung einer Long List an makroökonomischen Indikatoren in eine Short List

In diesem Kapitel soll die Überleitung einer Long List an makroökonomischen Daten anhand statistischer Kriterien in eine Short List erläutert werden. Diese selektierte Short List findet Eingang in die Naive Bayes Analysen und beeinflusst somit in weiterer Folge direkt die Entscheidung, ob Anleger in einen Index investieren sollten oder ob sie besser beraten wären, ihr Geld vom Markt abzuziehen. Bevor auf die im Selektionsprozess angewendeten statistischen Kriterien und Verfahren in jener Reihenfolge, in welcher diese zum Einsatz gekommen sind, detaillierter eingegangen wird, wird ein Überblick über lineare Regressionen im Allgemeinen sowie multiple lineare Regressionen im

Besonderen gegeben und auf die Besonderheiten und Herausforderungen in der Arbeit mit Zeitreihen eingegangen.

4.1 Lineare Regressionen

Um jene makroökonomischen Indikatoren zu identifizieren, die signifikant dazu beitragen, die Entwicklung des jeweiligen Aktienindex zu erklären und in weiterer Folge als Basis für Investitionsentscheidungen beim „Naive Bayes Investing“ dienen, wurden mit sämtlichen Indikatoren, deren Einfluss im Rahmen einer vorangegangenen Literaturrecherche als hoch oder sehr hoch klassifiziert wurde sowie gewisse statistische Erfordernisse (siehe Abschnitt 4.4 bis 4.6) erfüllen, multiple lineare Regressionsgleichungen aufgestellt. Ziel von Regressionsgleichungen ist es, zu eruieren, inwieweit eine oder mehrere Einflussvariablen das Eintreten einer Zielvariablen erklären können. Gegeben der Annahme, dass die erklärende Variable die abhängige linear beschreibt, wird der Zusammenhang durch eine lineare mathematische Funktion beschrieben (Frost, 2017, S. 5). „Der Zusammenhang ist nicht deterministisch sondern durch zufällige Fehler additiv überlagert“ (Frost, 2017, S. 5).

$$\hat{Y}_t = \alpha + \beta X_t + \varepsilon_t \quad (31)$$

$\hat{Y}_t$	abhängige, erklärte Variable
α	Achsenabschnitt
X_t	Regressor (unabhängige, erklärende Variable)
β	Regressionskoeffizient für die erklärende Variable
ε_t	Störgröße (diese überlagert den linearen Zusammenhang)

Die lineare Abhängigkeit wird überlagert vom Einfluss der Störgröße ε , wobei unterstellt wird, dass die Störgröße zufällig ist und ihr Erwartungswert 0 entspricht (Schuster, 2017, S. 219).

Das populärste Instrument für Modellierungen dieser Art ist die Methode der kleinsten gemeinsamen Quadrate („Method of least Squares“).

Zur Vermeidung von Ergebnisverzerrungen werden die Indikatoren auf Autokorrelation getestet und gegebenenfalls statt den Basis-Indikatoren deren Differenzen in die Regressionsgleichung aufgenommen. Danach wird mittels VIF-Test auf Vorliegen von Multikollinearität getestet und die Indikatoren-Auswahl mittels „Forward Selection“- und

„Backward Elimination“-Methode auf jene mit entsprechender Grenzerklärungskraft eingeschränkt.

4.2 Multiple lineare Regressionen

Entsprechend der einfachen linearen Regression gilt es auch bei der multiplen linearen Regression die quadratische Abweichung zwischen tatsächlichem Wert Y_t , der beobachtet wurde, und dem Schätzwert $\hat{Y}_t$ auf der Regressionsfläche zu minimieren (Schuster, 2017, S. 242). Im Gegensatz zur einfachen linearen Regression, bei der davon ausgegangen wird, dass eine vorliegende Beobachtung insbesondere von einem bestimmten Merkmal abhängig ist und sämtliche andere Einflüsse, denen diese Beobachtung unterliegt, nur unwesentlich oder zufällig sind, setzt sich die multiple lineare Regression zum Ziel einen linearen Zusammenhang zwischen mehreren unabhängigen Einflussgrößen und der endogenen Variablen herzuleiten (Schuster, 2017, S. 219). Die nachfolgende Gleichung stellt die Grundform eines multiplen linearen Regressionsmodells dar:

$$\hat{Y}_t = \alpha + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_p X_p + \varepsilon_t \quad (32)$$

$\hat{Y}_t$	abhängige, erklärte Variable
α	Achsenabschnitt
$X_1, X_2, \dots, X_p$	Regressoren (unabhängige, erklärende Variablen)
$\beta_1, \beta_2, \dots, \beta_p$	Regressionskoeffizienten für die erklärenden Variablen
ε_t	Störgröße (diese überlagert den linearen Zusammenhang)

Während der Achsenabschnitt den Wert der abhängigen Variablen angibt, wenn die Ausprägung aller unabhängigen Variablen mit 0 angenommen wird, messen die Regressionskoeffizienten β_i die Veränderung der abhängigen Variablen Y bei Änderung der jeweiligen erklärenden Variable X um eine Einheit bei konstantem Niveau der übrigen Regressoren. Die Variable β ist somit der Grenzeffekt auf die erklärte Variable bei Adaptierung von X . ε steht für die Störgröße im Modell, welche den linearen Zusammenhang überlagert. Hinsichtlich der Störterme wird bei multiplen Regressionsmodellen unterstellt, dass diese normalverteilt sind, Homoskedastizität vorliegt (gleiche Varianz der Störterme) und die „assumption of no autocorrelation“ (Unabhängigkeit der Ausprägung von ε zu einem bestimmten Zeitpunkt von zeitlich vor- und nachgelagerten Werten) gilt (Hatekar, 2010, S. 272-274).

Weil es die multiple lineare Regression ermöglicht, die marginale Auswirkung auf die Zielvariable, die einzig dem jeweiligen Regressor zugeschrieben wird, für sämtliche erklärende Variablen des Modells zu schätzen, findet diese vielfach Gebrauch in der angewandten Forschung (Hatekar, 2010, S. 272f).

4.3 Herausforderungen in der statistischen Analyse von Zeitreihen

Während sich Querschnittsdaten auf zeitlich ungeordnete Beobachtungen für unterschiedliche Merkmalsträger, die meist zum gleichen Zeitpunkt erhoben werden, beziehen, handelt es sich bei Zeitreihendaten, um Daten, die in regelmäßigen Intervallen erhoben werden. Die Verwendung von Zeitreihendaten als erklärende Variablen in einer (multiplen) Regressionsanalyse bringt – im Gegensatz zu Querschnittsdaten – zusätzliche Spezifika mit sich, die es entsprechend zu adressieren gilt (Koop, 2005, S. 10). Zum einen können Zeitreihendaten einander nicht nur unmittelbar sondern häufig erst mit einem gewissen „time lag“ beeinflussen, zum anderen treten bei nicht-stationären Variablen häufig so genannte Scheinregressionen („spurious regression“) auf (Koop, 2005, S. 121).

Die erwähnten Besonderheiten bei Zeitreihen machen den Einsatz zusätzlicher statistischer Analysen beziehungsweise Verfahren notwendig, da andernfalls die durchgeführten Regressionen wenig aussagekräftig oder aber auch schlichtweg falsch sein können (Koop, 2005, S. 101).

Die Thematik betreffend den zeitlich verzögerten Einfluss durch Zeitreihen-Variablen wird anhand des Beispiels der Anpassung der Leitzinssätze durch die Nationalbank erläutert. Soll etwa eine stagnierende Wirtschaft durch eine Senkung der Leitzinsen angekurbelt werden, so dauert dies meist bis zu einem Jahr, bis sich diese Adaptierung in der gesamten Wirtschaft bemerkbar macht. Im Rahmen einer Regressionsanalyse wird der Wert der abhängigen Variablen zu einem bestimmten Zeitpunkt somit nicht allein über den Wert der unabhängigen Variablen im selben Zeitpunkt, sondern auch über deren Werte in vorangegangenen Zeitpunkten erklärt. In diesem Zusammenhang spricht man von einem „Distributed Lag“ Regressionsmodell der Form:

$$\hat{Y}_t = \alpha + \beta_0 X_t + \beta_1 X_{t-1} + \dots + \beta_q X_{t-q} + \varepsilon_t \quad (33)$$

Die erklärenden Variablen werden in diesem Fall als „Lagged Variablen“ bezeichnet, der Index q gibt die „Lag Length“ oder „Lag Order“, die maximale Anzahl an zeitlich verzögerten Perioden, die Eingang in das Regressionsmodell finden, an (Koop, 2005, S. 121f). Die zur Ausgangsvariable zugehörigen zeitlich-verzögerten Periodenwerte sind dabei als eigenständige erklärende Variablen zu betrachten. A-priori ist in den meisten Fällen nicht eindeutig, wie viele Lags im Modell inkludiert werden sollten. Eine methodisch einfache, wenn auch rechnerisch aufwendige Möglichkeit bietet hierbei der Einsatz von t-Tests. Man beginnt zunächst mit einer relativ großen Anzahl an Lags und testet, ob der p-Wert für die Hypothese $\beta_q = 0$ größer dem gewählten Signifikanzniveau ist. Trifft dies zu, so wird die Lag-Länge um eine Periode gekürzt und erneut getestet. Diese Vorgehensweise wird so lange wiederholt, bis die Hypothese verworfen werden kann. Die Anzahl an Beobachtungen in einem „Distributed-Lag-Modell“ entspricht schlussendlich der ursprünglichen Anzahl an Datenpunkten abzüglich der Lag-Länge (Koop, 2005, S. 129). Sollen für die multiple Regression mehrere erklärende Variablen herangezogen werden, so müssen alle Variablen dieselbe Anzahl an Datenpunkten aufweisen (Koop, 2005, S. 123).

4.4 Augmented Dickey Fuller Test und Differenzbildung

Für Zeitreihendaten gelten die vier Komponenten Trend, Zyklen, saisonale Effekte und unregelmäßige Schwankungen als charakteristisch, wobei nicht automatisch jede Zeitreihe all diese Elemente aufweisen muss, um als eine solche zu gelten. Während der Zyklus die „Hochs“ und „Tiefs“ der Datenreihe mit einer Periodenlänge von mehr als einem Jahr beschreibt, beziehen sich Saisonkomponenten auf periodische Schwankungen innerhalb eines Jahres (etwa monats- oder quartalsweise Zyklen) (Black, 2019, S. 913f). Beispiele hierfür sind erhöhte Konsumausgaben in der Weihnachtszeit und eine rege Jobnachfrage im Sommer (z.B. Studierende auf der Suche nach Praktika) und lassen keine Rückschlüsse über grundlegende Änderungen einer Volkswirtschaft zu (Baumohl, 2012, S. 25). Unregelmäßige beziehungsweise irreguläre Schwankungen hingegen finden in noch kürzeren Zeitspannen als Saisonkomponenten statt und umfassen Ausreißer, die durch historische Ereignisse bedingt sein können aber auch

häufig nicht genauer erklärt werden können. Der Trend wiederum gibt die langfristige Grundrichtung der Daten an (Aufwärts- oder Abwärtstrend) (Black, 2019, S. 913f).

Zur Bereinigung von Datenerhebungen um saisonale Faktoren („Seasonal Adjustments“), verwenden Experten historische Daten, die in der Regel 5 bis 10 Jahre in der Vergangenheit liegen, um wiederkehrende Muster (Saisonalitäten) zu erkennen und die Indikatoren um diese Effekte zu bereinigen („Seasonal Adjustments“). Da die in der Arbeit verwendeten Indikatoren ohnehin bereits um saisonale Komponenten bereinigt wurden, wird auf die Saisonkomponenten von Zeitreihendaten nicht weiter eingegangen (Baumohl, 2012, S. 26).

Hinsichtlich Trends lassen sich deterministische von stochastischen unterscheiden. Bei deterministischen Trends ist die Zeitreihe als Funktion der Zeit samt „white noise“ definiert (Gehrke, 2019, S. 231). In diesem Zusammenhang spricht man von „white noise“ oder „weißem Rauschen“, wenn die Zufallsvariablen unkorreliert sind (Wolf und Best, 2010, S. 1057). Folgt die Zeitreihe hingegen einem stochastischen Trend, hat diese einen Erwartungswert, der zeitunabhängig ist, wie dies etwa beim „Random Walk“ der Fall ist (Wolf and Best, 2010, S. 1061). Bei Vorliegen eines stochastischen Trends ist die Zeitreihe somit rein vom Zufall abhängig und kann daher auch nicht exakt bestimmt werden (Wolf, 2009, S. 455).

Zu den Zeitreihen mit stochastischer Trendausprägung zählen viele ökonomische Zeitreihen, wie etwa das BIP, der private Verbrauch oder Produktionsindex. Diese werden durch einen stochastischen Trend folgender Form beschrieben (Winker, 2017, S. 267):

$$\hat{Y}_t = \mu + Y_{t-1} + \varepsilon_t \quad (34)$$

Die beschriebene Größe $\hat{Y}$ im Zeitpunkt t hängt von dessen Niveau der Vorperiode (Y_{t-1}) sowie dem Zuwachs μ in jeder Periode und der unabhängig normalverteilten Störgröße ε_t ab. Diese Ausprägung des stochastischen Trends wird auch als Random Walk mit Drift bezeichnet, wohingegen man von einem einfachen Random Walk spricht, wenn $\mu = 0$ gilt (Winker, 2017, S. 267). Im Gegensatz zu deterministischen Trends haben stochastische ein langes Gedächtnis, „da jeder Schock ε_t , der in Y_t eingeht, auch für alle folgenden Perioden Relevanz hat“ (Winker, 2017, S. 267). Stochastische Trends weisen aufgrund

dieses langen Gedächtnisses nach Schocks meist längerfristige Abweichungen als deterministische Trends, die rasch wieder zur Trendlinie zurückkehren, auf (Winker, 2017, S. 268).

Die Differenzierung zwischen deterministischem und stochastischem Trend ist deshalb von Bedeutung, da Zeitreihen mit deterministischem Trend gemeinhin als stationär gelten, wohingegen solche mit stochastischem Trend nichtstationär sind (Winker, 2017, S. 267). „Im Hinblick auf die Stationaritätseigenschaft sind der Mittelwert und die Varianz einer Zeitreihe von besonderer Bedeutung. Wenn diese über den gesamten Zeitraum konstant sind, liegt Stationarität vor. Stationarität bedeutet demnach, dass die einzelnen Beobachtungen einer Zeitreihe um einen festen Mittelwert fluktuieren. Nichtstationäre Zeitreihen unterliegen einer Trendentwicklung und haben deshalb keinen festen Mittelwert“ (Wolf, 2009, S. 454f). Nichtstationäre Variablen können die Analyse von Zeitreihendaten und somit die Modellaussage verfälschen und dürfen daher in Regressionen nicht berücksichtigt werden. Da viele ökonomische Zeitreihen ein Langfristgedächtnis aufweisen und sich aus der damit einhergehenden Nichtstationarität Auswirkungen auf die Ergebnisse von Regressionsanalysen sowie deren Interpretation ergeben, gilt es, der Unterscheidung von deterministischem und stochastischem Trend besondere Aufmerksamkeit zu schenken (Winker, 2017, S. 267).

Zur Testung, ob eine Zeitreihe das Stationaritätskriterium erfüllt (so genannte Unit-Root-Tests), wird häufig der von D. Dickey und W. Fuller entwickelte einfache Dickey-Fuller-Test oder der Augmented-Dickey-Fuller-Test angewendet (Wolf, 2009, S. 454). Der einfache Dickey-Fuller-Test wurde im Jahr 1979 von Dickey und Fuller entwickelt und ist vom Ansatz der einfachste Test zur Überprüfung der Nullhypothese, dass die analysierte Zeitreihe nichtstationär ist. Zu deren Testung wird die Dickey-Fuller-Teststatistik herangezogen, die der bekannten t-Statistik entspricht. Bei der Testung der Nullhypothese weist die t-Statistik dann nicht die übliche t-Verteilung auf, wenn die Nullhypothese der Nichtstationarität zutrifft. In diesen Fällen müssen zur Beurteilung die kritischen Werte der Dickey-Fuller Verteilung herangezogen werden. Diese sind in der Regel erheblich größer als für die t-Verteilung. Abhängig davon, welche deterministischen

Konstanten (mit Konstante, mit Konstante und Trend, ohne Konstante) berücksichtigt werden, sind verschiedene kritische Werte erheblich (Winker, 2017, S. 273).

Der einfache Dickey-Fuller-Test kann allerdings nicht verwendet werden, wenn die dahinterliegende Annahme, dass Störterme unabhängig verteilt, also nicht autokorreliert sind, nicht hält. In derartigen Fällen werden Überprüfungen auf Nichtstationarität mittels Augmented-Dickey-Fuller-Test (in weiterer Folge als „ADF-Test“ abgekürzt) durchgeführt. Dies ist insbesondere dann der Fall, wenn die Struktur der Zeitreihe mehr als einen verzögerten Wert aufweist. Beim ADF-Test wird die Schätzgleichung des Dickey-Fuller-Tests um verzögerte Differenzen von Y_t erweitert. Für die Bestimmung der Anzahl an verzögerten Differenzen beziehungsweise die sogenannte Lag-Länge wird in der Praxis das Informationskriterium herangezogen (Winker, 2017, S. 275f). Informationskriterien beheben das Spannungsfeld, das sich aus dem Einbeziehen zusätzlicher verzögerter Terme ergibt: Dem steigenden Erklärungsgehalt steht eine Zunahme des Kollinearitätsproblems gegenüber. Die Kriterien, die in der Praxis am häufigsten eingesetzt werden, sind jene, die von Schwarz, Akaike sowie Hannan und Quinn entwickelt wurden. All diese Informationskriterien haben gemeinsam, dass die Varianz der geschätzten Residuen möglichst geringgehalten werden soll, da dies gleichbedeutend mit einer hohen Erklärungskraft des Modells ist. Da aber mit einer steigenden Anzahl an unabhängigen Variablen automatisch eine steigende Erklärungskraft des Modells einhergeht, werden „Strafterme“ in den Informationskriterien berücksichtigt. „Strafterme“ tragen diesem Umstand Rechnung, indem mit steigender Anzahl der Erklärungsvariablen das Informationskriterium wiederum erhöht wird (Winker, 2017, S. 251). „Die praktische Umsetzung auf Basis eines ausgewählten Informationskriteriums sieht so aus, dass zunächst eine Obergrenze p für die Lag-Länge festgelegt werden muss. Dann werden die Modelle mit $K = 0, \dots, p$ verzögerten Werten geschätzt und jeweils der Wert des Informationskriteriums berechnet. Ausgewählt wird die Spezifikation, d.h. der Wert von K , für den das Informationskriterium seinen minimalen Wert erreicht“ (Winker, 2017, S. 252).

Aufgrund der erwähnten Problematiken, die mit nichtstationären Zeitreihen einhergehen, müssen diese entweder aus Regressionsmodellen ausgeschlossen oder in eine

stationäre Zeitreihe überführt werden. Die Überführung einer nichtstationären stochastischen Trendreihe in eine stationäre, erfolgt mittels Differenzbildung ($\Delta X_t = X_t - X_{t-1}$). Weil die Differenzbildung eine nichtstationäre Zeitreihe in eine stationäre überführen kann, werden diese als differenzstationär bezeichnet. Mit Hilfe des ADF-Tests wird dann festgestellt, ob die Bildung der ersten Differenz tatsächlich zu einer stationären Zeitreihe geführt hat (Wolf, 2009, S. 455). Die erste Differenz ist mathematisch als die Veränderung beziehungsweise das Wachstum der betrachteten Variablen anzusehen (Koop, 2005, S. 136). Obwohl die Bildung der ersten Differenz für eine Vielzahl von ökonomischen Variablen die geforderte Stationarität mit sich bringt, zeigen Indikatoren, wie etwa das Preisniveau, dass auch Zeitreihen, die auf der Bildung der ersten Differenz basieren, nicht unbedingt das Stationaritätskriterium erfüllen müssen. Ist dies der Fall, so muss eine weitere Differenz gebildet werden (Dreger et al., 2013, S. 210).

Neben der Überführung von nichtstationären in stationäre Zeitreihen werden mit der Differenzbildung auch die in Zeitreihen häufig vorkommenden Problematiken Autokorrelation und Scheinregression behoben.

Autokorrelationen beschreiben Korrelationen innerhalb eines Beobachtungskriteriums. Als Beispiel sei hier etwa das BIP genannt, das in der Regel eine starke Korrelation mit dem BIP eines vorangegangenen Erhebungszeitpunkts aufweist. Das Vorliegen von Autokorrelation ist vor allem bei makroökonomischen Daten naheliegend, da sich diese im Zeitverlauf – und selbst in Krisen – nur langsam verändern. Aktuelle Messungen eines makroökonomischen Indikators weisen somit nicht nur eine starke Korrelation mit Vorperioden sondern auch mit „Y lagged p periods“ auf (Koop, 2005, S. 138f).

Zu den weiteren Problematiken, die durch nichtstationäre Zeitreihen bedingt sind, zählen auch Scheinregressionen. Von einer Scheinregression („spurious regression“) spricht man, wenn die Regressionsanalyse Zusammenhänge zwischen Variablen als statistisch signifikant ausweist, die in der Tat aber völlig unabhängig voneinander sind. Das wohl bekannteste Beispiel für Scheinregressionen sind der fälschlich postulierte Zusammenhang zwischen der menschlichen Geburtenrate und der Zahl der Storchpaare (Winker, 2017, S. 269f).

Zusammenfassend sei festgehalten, dass eine Vielzahl an makroökonomischen Zeitreihen einen stochastischen Trend aufweisen, die das erforderliche Stationaritätskriterium nicht erfüllen. Mittels Differenzbildung gelingt nicht nur die Überführung in eine differenzstationäre Zeitreihe, sondern es werden auch für nichtstationäre Zeitreihen typische Problematiken wie etwa Autokorrelation und Scheinregression behoben (Koop, 2005, S. 138).

4.5 Multikollinearität und VIF-Test

Nach einer Selektion der makroökonomischen Indikatoren anhand der eben erläuterten statistischen Kriterien wurden diese Variablen auf das Vorliegen von Multikollinearität überprüft.

Bei der Anwendung von Regressionsmodellen wird unterstellt, dass die erklärenden Modellvariablen nicht linear voneinander abhängen. Wird die Grundannahme verletzt, so liegt Multikollinearität vor, eine gewisse Form der linearen Abhängigkeit zwischen zwei oder mehreren erklärenden Variablen in einer multiplen Regressionsgleichung (Wolf, 2009, S. 221). Multikollinearität ist in erster Linie ein Phänomen in den Wirtschaftswissenschaften, da meist sowohl abhängige als auch unabhängige Variablen beobachtet und nicht im Rahmen von Experimenten gemessen werden (Dreger et al., 2013, S. 66). „Dadurch wird eine Messung des isolierten Einflusses einer Einflussgröße auf die zu erklärende Variable erschwert, da die exogenen Variablen häufig miteinander korrelieren“ (Dreger et al., 2013, S. 66).

Die Ursachen für Multikollinearität sind vielfältiger Natur: Prinzipiell gilt, dass die Wahrscheinlichkeit, dass zumindest zwei unabhängige Variablen eines Regressionsmodells miteinander korrelieren, zunimmt, wenn zusätzliche Erklärungsvariablen ins Modell aufgenommen werden. Aber auch ein paralleler Datenverlauf sowie Lag- und Carry-Over-Effekte bei Zeitreihenanalysen können zu derartigem Phänomen führen. In den meisten Fällen sind die Ursachen für Multikollinearität allerdings nicht auf den ersten Blick ersichtlich – erst eine detaillierte Datenanalyse ermöglicht eine Aufdeckung (Wolf, 2009, S. 222f).

Ist Multikollinearität gegeben, wird der Einfluss einer exogenen auf die endogene Variable zumindest teilweise durch eine weitere Erklärungsvariable beeinflusst. Es ist somit schwierig festzustellen, welche der erklärenden Variablen, in welchem Ausmaß Einfluss auf die abhängige Variable nimmt. Aus diesem Grund sind bei einer multiplen linearen Regression die Regressionskoeffizienten der korrelierten erklärenden Variablen geringer als diese bei der Durchführung von einfachen linearen Regressionsmodellen mit selbigen Variablen ausfallen. Der Gesamteffekt dieser Variablen auf die abhängige Variable wird allerdings in beiden Fällen korrekt erfasst (Dreger et al., 2013, S. 66). Obwohl diese erklärenden Variablen gemeinsam eine große Erklärungskraft für das aufgestellte Modell aufweisen, macht es deren Multikollinearität der Regression unmöglich deren Einzeleinfluss korrekt abzubilden. Ein wichtiges Indiz für das Vorliegen von Multikollinearität kann eine niedrige t-Statistik und ein daher hoher p-Wert bilden (Koop, 2005, S. 100). Darüberhinausgehend ist auch die Aussagekraft des Regressionsmodells eingeschränkt, da die Regressionskoeffizienten deutlich über- oder unterschätzt werden können. Weiters können durch Multikollinearität bedingte hohe Varianzen oder Standardfehler dazu führen, dass Koeffizienten Vorzeichen aufweisen, die im Widerspruch zur allgemeinen Logik stehen (Wolf, 2009, S. 221). In der Regel wird diesen mit Multikollinearität einhergehenden Problemen begegnet, indem diese Variablen aus dem Modell ausgeschlossen werden (Koop, 2005, S. 101).

Zu den konventionellen Methoden zur Aufdeckung von Multikollinearität zählt unter anderem die Korrelationsmatrix. Mittels Korrelationsmatrix können jeweils zwei unabhängige Variablen einer Regression auf das Vorliegen von Multikollinearität überprüft werden. Eine Korrelationsmatrix ist allerdings nur im Rahmen von bivariaten Regressionsanalysen eine ausreichende Methode, da im multivariaten Fall aufgrund von kombinierten Effekten von mehreren Variablen selbst dann Multikollinearität vorliegen kann, wenn die bivariaten Korrelationskoeffizienten keine derartigen Rückschlüsse zulassen würden. Während zwischen metrischen Variablen der Korrelationskoeffizient nach Bravais-Pearson herangezogen wird, wird zwischen ordinal skalierten Merkmalen der Rangkoeffizient nach Spearman ermittelt. Die Festlegung eines Signifikanzniveaus ist für die Prüfung auf Multikollinearität nicht zwangsweise erforderlich, da selbst eine nicht-signifikante hohe Korrelation einen Hinweis auf Multikollinearität darstellen kann

(Wolf, 2009, S. 224). Ein Korrelationskoeffizient kann einen Wert zwischen 1 und -1 annehmen. Da aber selbst bei geringen Koeffizienten Multikollinearität vorliegen kann, dient der Korrelationskoeffizient gemeinhin nur als Ausgangspunkt für detailliertere Analysen.

Die Methode der Hilfsregressionen bietet detaillierteren Aufschluss bei der Überprüfung auf Multikollinearität, ist allerdings mit einem erhöhten rechnerischen Aufwand verbunden. Bei diesem Verfahren wird für jede exogene Variable eine Regression auf die übrigen erklärenden Modellvariablen durchgeführt (Dreger et al., 2013, S. 76). Nimmt der multiple Korrelationskoeffizient R_i^2 einen Wert nahe 1 an, so deutet dies auf Multikollinearität hin, wobei derartige Probleme bereits auch bei einem wesentlich geringeren R_i^2 gegeben sein können. Der Toleranzwert stellt mit $TOL_i = 1 - R_i^2$ das Komplement des multiplen Korrelationskoeffizienten dar. Prinzipiell gilt daher, dass kleinere Werte das Vorliegen von Multikollinearität indizieren (Wolf, 2009, S. 225).

Der Variance-Inflation-Factor (in weiterer Folge als „VIF-Faktor“ abgekürzt) ist eine weitere Kennzahl, die auf Hilfsregressionen basiert und zur Aufdeckung von Multikollinearität eingesetzt wird. Diese Kennzahl gibt die Erhöhung der Varianz an, die sich dadurch ergibt, dass die getestete exogene Variable mit den übrigen Erklärungsvariablen entsprechend dem Bestimmtheitsmaß R_i^2 korreliert. Der Toleranzwert bildet demnach den Residualwert an der Gesamtvarianz, der nicht über die restlichen unabhängigen Variablen erklärt werden kann. Das Bestimmtheitsmaß ist somit Ausdruck jenes Anteils an der Gesamtvarianz einer bestimmten unabhängigen Variablen, der durch die restlichen exogenen Modellvariablen erklärt wird. Berechnet wird der VIF-Faktor auf Grundlage von Korrelationsmatrizen unabhängiger Variablen. Der VIF-Faktor ist der Kehrwert des Toleranzwerts und wird folgendermaßen ermittelt (Wolf, 2009, S. 225):

$$VIF = \frac{1}{1 - R_i^2} \quad (35)$$

Entsprechend den obenstehenden Ausführungen gilt auch für den VIF-Wert, dass höhere Werte eher Rückschlüsse auf das Vorhandensein von Multikollinearität erlauben. Als kritische Richtgröße gilt allgemein ein Wert von ≥ 10 . Das Vorliegen von Multikollinearität

kann aber auch bei geringeren Werten ohne die Durchführung weiterer Tests und Analysen nicht mit Sicherheit ausgeschlossen werden (Wolf, 2009, S. 225).

In der Praxis wird zur Verringerung von Multikollinearität in Modellen häufig darauf zurückgegriffen, jene Werte auszuschließen, deren VIF-Faktor den Schwellwert von 10 überschreitet. Es wird davon ausgegangen, dass deren Grenzerklärungskraft für das Gesamtmodell aufgrund ihrer hohen Korrelation mit einer anderen exogenen Variablen ohnehin nur gering ist. Derartige Regressoren können somit ausgeschlossen werden, ohne den Erklärungswert des Modells erheblich zu verringern (Wolf, 2009, S. 229).

4.6 Schrittweise Selektion der Variablen

Nachdem jene Variablen aus dem Modell entfernt wurden, die aufgrund von Multikollinearität zu einer verzerrten Aussage führen könnten, erfolgt die Selektion der unabhängigen Variablen in einem schrittweisen Auswahlprozess. Schrittweise Verfahren kommen insbesondere dann zum Einsatz, wenn sich die Beschränkung der unabhängigen Variablen weder aus der Forschungsfrage noch durch bereits durchgeführte Analysen ergibt – so wie dies in der vorliegenden Arbeit der Fall ist. Für den schrittweisen Selektionsprozess werden im Rahmen dieser Arbeit die Verfahren „Forward Selection“ und „Backward Elimination“ angewendet.

Die Forward Selection startet zunächst mit 0 erklärenden Variablen (Anderson, 2020, S. 759). In weiterer Folge wird dann jene Variable ermittelt, die den höchsten t-Wert aufweist, der das Signifikanzkriterium erfüllt, und deren Erklärungswert für die abhängige Variable am größten ist, also das höchste R^2 aufweist. Wurde eine Variable im Auswahlverfahren der Forward Selection einmal ins Modell aufgenommen, so ist ein Ausschluss der gleichen Variablen im weiteren Verlauf ausgeschlossen. Mögliche Korrelationen zwischen den unabhängigen Modellvariablen bleiben somit unberücksichtigt. Im nächsten Schritt wird eine erneute Analyse der verbliebenen Variablen durchgeführt und wiederum jene mit dem höchsten R^2 beziehungsweise t-Wert im Modell inkludiert (Black, 2019, S. 873). Dieses Verfahren kann so oft wiederholt werden bis alle Regressorvariablen ins Modell aufgenommen wurden, oder es wird dann abgebrochen, wenn keine der verbliebenen Variablen mehr einen vordefinierten p-Wert (p-Value-to-enter) erreicht und somit das gewählte Signifikanzniveau verfehlt (Anderson,

2020, S. 759). Es sei an dieser Stelle darauf hingewiesen, dass der Rang, mit dem eine unabhängige Variable ins Regressionsmodell aufgenommen wird, prinzipiell keine Aussage betreffend deren Bedeutung für die Erklärung der abhängigen Variablen zulässt. Dies gilt nicht für jene Variable, die als erstes aufgenommen wird, allerdings für alle darauffolgenden. Deren Bedeutung für das Modell bemisst sich nämlich nicht am Erklärungswert, den der betreffende Regressor alleine für die abhängige Variable hat sondern daran, wie viel von der Restabweichung (residual variation) durch die Aufnahme dieser Variablen ins Modell erklärt werden kann (Black, 2019, S. 894).

Das Verfahren der Backward Elimination ist in seinen Grundzügen jenem der Forward Selection sehr ähnlich, wählt allerdings einen anderen Ansatzpunkt. Während Ausgangspunkt der Forward Selection ein „leeres Regressionsmodell“ bildet, werden bei der Backward Elimination zunächst alle erklärenden Variablen in die Regressionsgleichung aufgenommen (Black, 2019, S. 873). Die Stichprobengröße gilt bei der Anwendung von Backward Elimination als Limiting Factor, da diese groß genug sein muss, um ein Ausgangsmodell zu rechtfertigen, das sämtliche vorselektierten Erklärungsvariablen umfasst. Ist die Anzahl der Datenpunkte nicht ausreichend groß, so ist Backward Elimination keine geeignete Methode um eine multiple Regressionsgleichung aufzustellen (Black, 2019, S. 894). Zunächst wird getestet, ob sich unter den Erklärungsvariablen solche befinden, die nicht signifikant sind. Sind sämtliche Variablen signifikant, so bleibt das Ausgangsmodell der Backward Elimination unverändert. Für den Fall, dass gewisse Variablen nicht signifikant sind, so wird zuerst die am wenigsten signifikante Variable ausgeschlossen und ein adaptiertes Modell gebildet. Die eben beschriebenen Schritte werden so lange wiederholt, bis sämtliche Regressorvariablen das gewählte Signifikanzniveau erfüllen, also keine Variable den p-Value-to-leave überschreitet (Black, 2019, S. 873).

In manchen Fällen führt die Anwendung der Backward Elimination zur selben Variablen-Auswahl wie die der Forward Selection, die Ergebnisse können gleichermaßen aber auch voneinander abweichen. Da bei Backward Elimination aufgrund des so genannten Suppressioneffekts die Wahrscheinlichkeit höher ist, unabhängige Variablen fälschlicherweise auszuschließen, obwohl diese zur Erklärung der abhängigen Variablen

beitragen würden, wird der Forward Selection generell der Vorzug gegeben (Field et al., 2012, S. 265). Von einem Suppressioneffekt spricht man dann, wenn eine Variablen, die zwar nur in geringem Ausmaß oder gar nicht mit der erklärenden Variable korreliert, aber „den Vorhersagebeitrag einer (oder mehrerer) anderer Variablen erhöht, indem sie irrelevante Varianzen in der (den) anderen Prädiktorvariablen unterdrückt“ (Bortz, 2005, S. 459).

5 Auswahl und Beschreibung relevanter Daten

Fokus dieses Kapitels ist neben einer prägnanten Erläuterung der Interdependenzen zwischen volkswirtschaftlichen Indikatoren und Aktienmärkten, der Auswahlprozess jener Indikatoren, die die Basis des „Naive Bayes Investing“ bilden. In diesem Zusammenhang wird zunächst auf die Interdependenzen zwischen makroökonomischen Indikatoren im Allgemeinen eingegangen, und es werden dann jene Indikatoren präsentiert, die potentiell für die Entwicklung der gewählten Vergleichsmärkte interessant sein könnten und daher einer eingehenden statistischen Analyse unterzogen werden.

In weiterer Folge wird auch auf die beiden Indizes „Deutscher Aktienindex“ sowie „Standard & Poor's 500“, die als Benchmark für die Auswertungen herangezogen werden, auf deskriptiver Ebene eingegangen.

5.1 Makroökonomische Indikatoren und deren Einfluss auf Aktienmärkte

In einer Vielzahl von Publikationen werden die Wechselwirkungen zwischen makroökonomischen Indikatoren und der Entwicklung von Aktienpreisen diskutiert. Darüberhinausgehend wird ebenso davon ausgegangen, dass die Aktienkurse zum Teil von fundamentalen volkswirtschaftlichen Variablen bestimmt werden (Pilinkus, 2010, S. 291).

Bevor auf die Relevanz makroökonomischer Indikatoren für den Equity Markt detaillierter eingegangen wird, erfolgt eine Begriffsverortung sowie die Darstellung verschiedener Klassifizierungsmöglichkeiten.

5.1.1 Definition und Klassifizierung

„Indikatoren (lat. indicare: anzeigen) sind Daten und Werte, die Auskunft über die Richtung und Intensität einer ökonomischen Variablen, zum Beispiel der Konjunktur, geben können“ (Wildmann, 2016, S. 84). Aus dieser Definition ergibt sich, dass Konjunkturindikatoren nicht nur als deskriptives Instrument Anwendung finden sondern auch zu Prognosezwecken zum Einsatz kommen. Konjunkturindikatoren können sowohl nach ihrer Vorlaufeigenschaft als auch nach ihrer Entstehungsart eingeteilt werden (Oppenländer, 1996, S. 23f).

Entsprechend ihrer Vorlaufeigenschaft kann zwischen vorlaufenden, gleichlaufenden und nachlaufenden Indikatoren unterschieden werden: Zu den Frühindikatoren, die Aufschluss über die künftige Entwicklung der Konjunktur geben, zählen etwa Investitionen, Baugenehmigungen sowie Auftragseingänge und -bestände. Präsenzindikatoren im Gegensatz dazu zeigen an, wie die aktuelle wirtschaftliche Lage ist. Als Beispiele für derartige Indikatoren seien an dieser Stelle ausgelastete Produktionskapazitäten sowie reduzierte Lagerbestände genannt. Das Preisniveau oder aber auch Gewinne und die Situation am Arbeitsmarkt wiederum zeigen retrospektiv das Konjunkturklima an und zählen demnach zu den Spätindikatoren, da sich etwa ein wirtschaftlicher Aufschwung erst mit einer gewissen zeitlichen Verzögerung in einem höheren Preis- und Gehaltsniveau niederschlägt (Wildmann, 2016, S. 85).

Hinsichtlich Art der Entstehung werden Konjunkturindikatoren in die Gruppen der quantitativen und qualitativen Indikatoren klassifiziert. Quantitative Indikatoren haben den Vorteil, dass diese numerisch und auch im Hinblick auf die statistische Messgenauigkeit den qualitativen Faktoren überlegen sind. Allerdings ist mit der Datenmanipulation im Zusammenhang mit der Erhebung von quantitativen Indikatoren häufig ein erheblicher Aufwand verbunden, so dass deren endgültigen Werte erst mit einiger Verzögerung publiziert werden. Insbesondere bei „leading indicators“, die – wie bereits angemerkt – als Vorboten der Konjunkturentwicklung dienen, ist diese Verzögerung problematisch, da der Vorteil des Frühboten somit zur Gänze ausgehebelt werden kann. Qualitative Indikatoren weisen demgegenüber den Benefit auf, dass diese rasch veröffentlicht werden können. Bei monatlich-publizierten Indikatoren liegen die finalen Daten meist

schon 14 Tage nach der letzten Datenerfassung vor. Damit die Daten auch von Repräsentativität zeugen, müssen zum einen eine Vielzahl an Datenpunkten erfasst werden, zum anderen muss die Antwortquote entsprechend hoch sein. Qualitative Indikatoren, die diese Anforderungen erfüllen, werden beispielsweise vom ifo-Institut (Institut für Wirtschaftsforschung) publiziert. Des Weiteren werden in bestimmten qualitativen Indikatoren die Erwartungen von Unternehmen erfasst und somit die Marktstimmung abgebildet, wobei im Gegensatz zu quantitativen Indikatoren keinerlei Trendbereinigung erforderlich ist (Oppenländer, 1996, S. 27f).

Ein weiteres Klassifikationskriterium von volkswirtschaftlichen Indikatoren ist die Differenzierung zwischen Einzel- und Gesamtindikatoren. Während Einzelindikatoren in der Regel bereits allein eine hohe Aussagekraft besitzen, ist vor allem für die Stabilität von „leading indicators“ das Zusammenspiel mehrerer – qualitativer als auch quantitativer – Indikatoren entscheidend. Da allerdings die Gewichtung, mit denen Einzelindikatoren in einen Gesamtindikator eingehen, bedeutsam ist und diese Gewichtung im Zeitverlauf veränderlich sein kann, sind historische Vergleiche meist wenig informativ, weshalb Gesamtindikatoren immer mehr ins Hintertreffen geraten sind (Oppenländer, 1996, S. 28).

5.1.2 Bedeutung makroökonomischer Indikatoren für den Aktienmarkt

Die Relevanz makroökonomischer Indikatoren für den Aktienmarkt soll durch folgenden Sachverhalt verdeutlicht werden: Vor rund 30 Jahren – als es in Amerika noch keinerlei Regelungen hinsichtlich der Publikation von Konjunkturberichten gab – ist es zu vielfachem Missbrauch im Zusammenhang mit der Publikation volkswirtschaftlicher Daten gekommen. So machten sich Politiker das Publikationsdatum bestimmter Indikatoren zu Nutze, um die Stimmung der Wählerschaft zu beeinflussen. Investmentfirmen wieder versuchten bereits im Vorfeld, Insider-Informationen über die Entwicklung bestimmter Indikatoren zu erhalten, um ihre Investments danach auszurichten. Um diesen Missbrauch zu unterbinden, wurden klare Vorschriften zur Veröffentlichung von Konjunkturindikatoren erlassen (Baumohl, 2012, S. 7).

Die wechselseitige Beziehung zwischen Aktienmarkt einerseits und Änderungen in der Ökonomie eines Landes andererseits sowie der Aktienmarkt als Spiegelbild der

gesamtwirtschaftlichen Situation eines Landes wurde auch vielfach in der Literatur diskutiert. So gilt, dass ein genereller Anstieg der Börse-Aktivität beobachtbar ist, sobald sich die wirtschaftliche Lage verbessert und ein anhaltender Preisverfall am Aktienmarkt wiederum eine bevorstehende Depression indizieren kann (Pilinkus, 2010, S. 292).

Obwohl die Interdependenzen zwischen Kapitalmarkt und Ökonomie nachgewiesen sind, sind sich Autoren über die Art und Richtungswirkung des Zusammenwirkens uneinig: "Literature on finance and development to a great extent suggests that there is a relationship between financial system development, macroeconomic variables and economic development but the exact nature of the relationship has been inconclusive. Economists hold divergent views on the nature of interrelationships between financial and economic development" (Kyereboah-Coleman und Agyire-Tettey, 2008, S. 368).

Volkswirtschaftliche wie fiskalpolitische Rahmenbedingungen sind Faktoren, die wesentlichen Einfluss auf die Profitabilität und den Gesamterfolg eines Unternehmens nehmen. Die Autoren halten aus der Vielzahl an makroökonomischen Indikatoren insbesondere Wechselkurse, Zinssätze und Inflationsraten als relevant für die Aktivitäten am Aktienmarkt, da diese direkt auf das Handeln von Unternehmen wirken (Kyereboah-Coleman und Agyire-Tettey, 2008, S. 369).

Die Selektion aus der Masse an veröffentlichten volkswirtschaftlichen Indikatoren, die letztlich Eingang in die jeweiligen Analysen finden, ist jedoch divergent. Pilinkus (2010, S. 294) schreibt, dass viele Autoren auf Indikatoren wie das BIP, Zinssätze und Indizes zur industriellen Produktion zurückgreifen. Häufig mangelt es aber einer fundierten Begründung, warum gerade diese Faktoren Berücksichtigung finden. Die genannten Indikatoren bilden meist nicht aufgrund ihrer bekannter Interdependenzen Ausgangspunkt weiterer Untersuchungen sondern eher wegen deren Popularität und Frequenz, mit welcher diese publiziert werden.

Die Höhe der Inflationsrate wirkt sich direkt auf das Investitionsverhalten von Unternehmen aus. Weisen die Inflationsraten etwa eine unvorhergesehene Höhe auf, so stehen Unternehmen oftmals – auch aufgrund der zunehmenden Unsicherheiten – vor der Entscheidung, geplante Investitionen zu verschieben. Instabilitäten des allgemeinen

Preisniveaus begünstigen zudem die Investition in Anlagen mit kurzem Zeithorizont, da längerfristige Prognosen kaum getroffen werden können. Ist das Wirtschaftsumfeld von erhöhten Inflationsraten geprägt, so reduzieren sich nicht nur die Anlagen in langfristige Investments, sondern es ist in der Regel auch ein Shift von Finanzanlagen hin zu materiellen Vermögenswerten beobachtbar. Ein Zusammenhang zwischen Inflationsrate und Aktienmarkt ergibt sich zudem daraus, dass mit einem Anstieg des allgemeinen Preisniveaus eine Zunahme der Lebenskosten einhergeht, weshalb Privatpersonen häufig dem Aktienmarkt Geld entziehen und dieses für Verbrauchsgüter ausgeben. Die reduzierte Nachfrage am Aktienmarkt führt zu einem geringeren Handelsvolumen und schlussendlich zu sinkenden Aktienpreisen. Aus den dargelegten Gründen wird ein negativer Zusammenhang zwischen Inflationsrate und Handelsvolumen sowie Aktienindex postuliert (Kyereboah-Coleman and Agyire-Tettey, 2008, S. 370f).

Zinssätze und Wechselkurse sind als Preis für Kredite beziehungsweise Fremdwährungen zu sehen und stehen direkt in Zusammenhang mit der Allokation von Ressourcen, Profitabilität und Produktionsniveau. Fluktuationen von Kreditzinsen und Wechselkursen schlagen sich daher auch in den Aktienpreisen nieder. Fallende Sparszinsen reduzieren die Attraktivität der Geldanlage in Sparbücher erheblich. Über einen Anstieg der Nachfrage nach Aktien wirken sich somit niedrigere Zinsen positiv auf das Preisniveau am Aktienmarkt aus. Steigen hingegen die Zinsen auf Staatsanleihen, tendieren Anleger dazu, verstärkt in staatliche Obligationen zu investieren. Da für Investoren mit hohem Grad an Risikoaversion Staatsanleihen häufig ein geeignetes Substitut für Aktienportfolios darstellen, führen steigende risikolose Zinssätze zu fallenden Aktienindizes. Anders verhält sich die Interdependenz zwischen Performance des Aktienindex und Höhe der Kreditzinsen: Aufgrund steigender Kreditzinsen nehmen Unternehmen zum Teil Abstand davon ihre Investitionen fremd zu finanzieren und beschaffen sich liquide Mittel stattdessen am Aktienmarkt. Zusätzliche Listings erhöhen die Handelsaktivitäten und können letztendlich einen Anstieg der Aktienkurse nach sich ziehen (Kyereboah-Coleman and Agyire-Tettey, 2008, S. 370).

5.2 Auswahlprozess und Erläuterung volkswirtschaftlicher Indikatoren

Die Auswahl von volkswirtschaftlichen Indikatoren, die als potentiell relevant für die Kurs-Entwicklungen am Aktienmarkt erachtet wurden, basiert auf einer Literaturrecherche, die neben renommierten Webseiten („Trading Economics“ (TRADING ECONOMICS, 2020) und „Investing.com“ (Fusion Media Limited, 2020)) auch das Werk von Bernhard Baumohl „The Secrets of Economic Indicators“ umfasst. Damit ein fraglicher Indikator für die weitere Analyse als wesentlich erachtet wurde, musste der genannte Autor die Auswirkung der Indikator-Entwicklung auf den Markt als „high“ beziehungsweise „very high“ einstufen oder die Analysten, die ihre Einschätzungen für die herangezogenen Webseiten treffen, die Relevanz des Indikators als „mittel“ oder „hoch“ (Deutscher Aktienindex) beziehungsweise „hoch“ (Standard & Poor's 500) einstufen.

Die den Analysen zugrunde gelegten volkswirtschaftlichen Daten sowie Aktienindizes wurden über den Finanzdatenanbieter Bloomberg (Bloomberg L.P., 2020) respektive Yahoo Finance (Finance.yahoo.com., 2020, 2021) bezogen.

5.2.1 Ökonomische Indikatoren mit potentieller Relevanz für die Vergleichsmärkte

Bundesregierungen aber auch Privatinstitutionen veröffentlichen auf täglicher, wöchentlicher, monatlicher oder quartalsweiser Basis Konjunkturberichte in mannigfaltiger Weise, wobei jeder dieser Berichte Rückschlüsse auf einzelne Segmente der Wirtschaft erlaubt: „By following these indicators, one can get the latest reading on the economy's health and valuable clues on where it is heading“ (Baumohl, 2012, S. 8).

Doch selbst die Auswertung und Interpretation sämtlicher verfügbarer Informationen über eine Volkswirtschaft erlaubt meist keine eindeutige Prognose über die Entwicklung des Marktes. Dies liegt unter anderem darin begründet, dass verschiedene Indikatoren widersprüchliche Signale über die voraussichtliche Entwicklung des Marktes liefern können. Zudem können einzelne Indikatoren in einem kontraintuitiven Zusammenhang zueinanderstehen. Als Beispiele sollen an dieser Stelle „Consumer Confidence“ und „Consumer Spending“ genannt werden. Während man erwarten würde, dass pessimistische Verbraucherstimmungen mit einem Rückgang der Ausgaben von Haushalten einhergehen, zeichnet sich in der Realität oftmals ein gänzlich anderes Bild. So konnten 2005 in Amerika trotz Hochphase der Konjunktur und damit einhergehender

optimistischer Verbraucherstimmung dennoch keine steigenden Konsumausgaben verzeichnet werden, da diese nicht allein von der Verbraucherstimmung sondern gleichermaßen von Faktoren wie Jobsicherheit, Entwicklung des persönlichen Vermögens und der Zinslage abhängen (Baumohl, 2012, S. 7f).

Für Anleger am Aktienmarkt sind insbesondere jene Indikatoren wesentlich, die Aufschluss über bevorstehende Veränderungen im Konsumverhalten von Haushalten und Unternehmen liefern können, da diese in weiterer Folge Auswirkung auf die Profitabilität und den Umsatz von börsennotierten Unternehmen haben. Des Weiteren gilt, dass Faktoren besonders relevant sind, falls diese mehrere der nachfolgend beschriebenen Eigenschaften aufweisen (Baumohl, 2012, S. 11f):

- **Genauigkeit:** Bestimmte Maßzahlen gelten aufgrund ihrer Datenerhebungsmethode, Rücklaufquote und ihres Stichprobenumfangs als besonders verlässlich.
- **Aktualität:** Indikatoren, die aktuell sind, sind für Investoren wesentlich relevanter. Etwa der Report zur Arbeitsmarktsituation ist für Investoren von besonderem Interesse, da dieser schon rund eine Woche nach Monatsende publiziert wird.
- **Phase des Konjunkturzyklus:** Abhängig davon, in welchem Stadium des Konjunkturzyklus sich die Wirtschaft gerade befindet, werden einige Indikatoren dringlicher erwartet als andere. Dementsprechend unterschiedlich wirkt sich die Publikation verschiedener Indikatoren auf den Aktienmarkt aus.
- **Vorhersagefähigkeit:** Zu den Indikatoren, denen am meisten Beachtung geschenkt wird, zählen jene, die bereits in der Vergangenheit als zuverlässiger Forecast für Wendepunkt in der gesamtwirtschaftlichen Situation eines Landes dienten (z.B.: „Durable Goods New Orders“, „Existing Home Sales“).

Wie bereits zuvor erwähnt, erfolgte die Vorauswahl von Indikatoren mit potentieller Relevanz auf Basis qualitativer Recherchen. Die in den folgenden Unterabschnitten dargestellten Indikatoren bildeten die Ausgangsbasis für den quantitativen Teil der vorliegenden Masterarbeit und wurden in weiterer Folge einer eingehenden statistischen Analyse unterzogen.

5.2.2 DAX ®

In der Tabelle 9 werden jene Indikatoren aufgelistet, deren Entwicklung potentiell Auswirkung auf den DAX haben könnte:

Long List Germany	Frequenz der Publikation
CPI All items MoM	Monatlich
ZEW Germany Assessment of current situation	Monatlich
ZEW Germany expectations of economic growth	Monatlich
Retail Sales constant prices SWDA MoM	Monatlich
Unemployment change	Monatlich
Unemployment rate	Monatlich
Ifo Pan germany Business Climate	Monatlich
Ifo pan germany Business Expectations	Monatlich
Ifo Pan germany Current Assessment	Monatlich
Germany Factory Orders (MoM)	Monatlich
Germany Industrial Production (MoM)	Monatlich
Germany producer prices MoM	Monatlich
Trade Balance EUR NSA	Monatlich
GfK Consumer Confidence	Monatlich
Manufacturing PMI	Monatlich
Services PMI	Monatlich
Exchange Rate EUR/USD	Täglich
EUR German Sovereign Curve 3M	Täglich
EUR German Sovereign Curve 6M	Täglich
EUR German Sovereign Curve 10Y	Täglich

Tabelle 9: Long List an makroökonomischen Indikatoren für den DAX

5.2.3 S&P 500 ®

Tabelle 10 zeigt jene Indikatoren, welche die eben beschriebenen Charakteristika aufweisen und aufgrund der in Abschnitt 5.2 genannten Selektionskriterien als potentiell relevant für die Entwicklung des S&P 500 und somit die weiteren Analysen der vorliegenden Masterarbeit eingestuft worden sind:

Long List U.S.	Frequenz der Publikation
ADP National Employment report SA Private Nonfarm Level Change	Monatlich
US Trade Balance of Goods and Services SA	Monatlich
US Employees on Nonfarm Payrolls Total MoM Net Change SA	Monatlich
US Unemployment Rate Total in Labour force Seasonally Adjusted	Monatlich
US Job Openings by industry total SA	Monatlich
US CPI Urban Consumers MoM SA	Monatlich
US CPI Urban Consumers Less Food & Energy MoM SA	Monatlich
US CPI Urban Consumers Less Food & Energy YoY SA	Monatlich
Adjusted Retail Sales Less Autos SA Monthly % Change	Monatlich
Adjusted Retail & Food Services Sales SA Total Monthly % Change	Monatlich
Private Housing Authorized by Building Permits by Type Total SAAR	Monatlich
US PPI Final Demand MoM SA	Monatlich
Philadelphia Fed Business Outlook Survey Diffusion Index General Conditions	Monatlich
US Existing Home Sales SAAR	Monatlich
ISM Manufacturing PMI SA	Monatlich
ISM Services PMI	Monatlich
CB Consumer Confidence SA	Monatlich
US Personal Income MoM SA	Monatlich
US Pending Home Sales Index MoM SA	Monatlich
Exchange Rate EUR/USD	Täglich
US Personal Consumption Expenditures Nominal Dollars MoM SA	Monatlich
US Treasury Actives Curve 3M	Täglich
US Treasury Actives Curve 6M	Täglich
US Treasury Actives Curve 10Y	Täglich
US Durable Goods New Orders Industries MoM SA	Monatlich
US Durable Goods New Orders Total ex Transportation MoM SA	Monatlich

Tabelle 10: Long List an makroökonomischen Indikatoren für den S&P 500

5.3 Charakteristika gewählter Benchmarks

Als Benchmark für die entwickelte Handelsstrategie dienen Aktienindizes: Definitionsgemäß ist ein Aktienindex „(...) eine Kennziffer zur Darstellung der Kursentwicklung oder Wertentwicklung (Performanceindex) von Aktien. Das Verhalten eines Aktienkursindex wird vor allem durch die Kurse der im Index enthaltenen Aktien beeinflusst, aber meist auch durch eine Gewichtung der Einzelwerte. Aktienindizes unterscheiden sich vor allem durch ihre Gestaltung (Kurs- oder Performanceindex), die Anzahl der enthaltenen Papiere oder durch die Index-Gewichtung“ (Boersenlexikon.faz.net., 2021). Gegenüber Einzelaktien weisen Aktienindizes den

erheblichen Vorteil auf, dass diese Investoren nicht nur Einschätzungen über die Gesamtsituation eines einzigen Unternehmens sondern über den gesamten für den jeweiligen Investor relevanten Markt oder Sektor erlauben (Pilinkus, 2010, S. 293).

Nachfolgend werden die beiden Indizes DAX, ein Performanceindex, und S&P 500, ein Kursindex, die als Benchmark für die Naive Bayes Handelsstrategie dienen sollen, erläutert. Die Entscheidung, die im Rahmen dieser Masterarbeit entwickelte Handelsstrategie an diesen Aktienindizes zu testen, basierte zum einen auf dem Streben Indizes großer Volkswirtschaften zu wählen, zum anderen sollte es hinsichtlich der gelisteten Unternehmen keine Überschneidungspunkte geben. Durch die Wahl eines europäischen und US-amerikanischen Börsenindex konnten die Handelsstrategien zudem – mit Deutschland als sozialer und USA als freier Marktwirtschaft – auf verschiedenen Märkten erprobt werden.

5.3.1 DAX ®

Der Deutsche Aktienindex (in weiterer Folge als „DAX“ abgekürzt) ist der Top-Index der Deutschen Börse und umfasst Aktienkurse jener 30 Unternehmen des Prime Standards der Frankfurter Wertpapierbörse, die am größten sind und die höchste Liquidität aufweisen („Blue Chips“). Der DAX dient als Basis für die Ableitung des Sub-Index „DAX ex Financials“, der sämtlichen Unternehmen, die in den Branchen Banken, Financial Services oder Immobilienwirtschaft tätig sind, exkludiert (Dax-indices.com., 2020a).

Die gesamte Marktkapitalisierung börsennotierter Aktiengesellschaften in Deutschland wird zu rund 80 % von den in der nachfolgenden Tabelle 11 dargestellten DAX-Unternehmen getragen (Dax-indices.com., 2020b):

Instrument	ISIN
ADIDAS AG	DE000A1EWWW0
ALLIANZ SE	DE0008404005
BASF SE	DE000BASF111
BAY.MOTOREN WERKE AG	DE0005190003
BAYER AG	DE000BAY0017
BEIERSDORF AG	DE0005200000
CONTINENTAL AG	DE0005439004
COVESTRO AG	DE0006062144
DAIMLER AG	DE0007100000
DELIVERY HERO SE	DE000A2E4K43
DEUTSCHE BANK AG	DE0005140008
DEUTSCHE BOERSE	DE0005810055
DEUTSCHE POST AG	DE0005552004
DEUTSCHE WOHNEN SE INH	DE000A0HN5C6
DT.TELEKOM AG	DE0005557508
E.ON SE	DE000ENAG999
FRESEN.MED.CARE KGAA	DE0005785802
FRESENIUS SE+CO.KGAA	DE0005785604
HEIDELBERGCEMENT AG	DE0006047004
HENKEL AG+CO.KGAA VZO	DE0006048432
INFINEON TECH.AG	DE0006231004
LINDE PLC EO 0,001	IE00BZ12WP82
MERCK KGAA	DE0006599905
MTU AERO ENGINES	DE000A0D9PT0
MUENCH.RUECKVERS.VNA	DE0008430026
RWE AG INH	DE0007037129
SAP SE	DE0007164600
SIEMENS AG	DE0007236101
VOLKSWAGEN AG VZO	DE0007664039
VONOVIA SE	DE000A1ML7J1

Stand 29.12.2020

Tabelle 11: Zusammensetzung des DAX

Betreffend den Anteil am Index ist „LINDE PLC“ (10,7 %) gefolgt von „SAP SE“ (10,5 %) und „SIEMENS AG“ (8,6 %) das größte Unternehmen. Werden die DAX-30-Unternehmen verschiedenen Sektoren zugeordnet, zeigt sich, dass der Großteil der Aktiengesellschaften in der Branche „Chemicals“ operiert. Die nachfolgende Übersicht zeigt eine gut diversifizierte Verteilung der DAX-Unternehmen auf verschiedene Branchen (Dax-indices.com., 2020c):

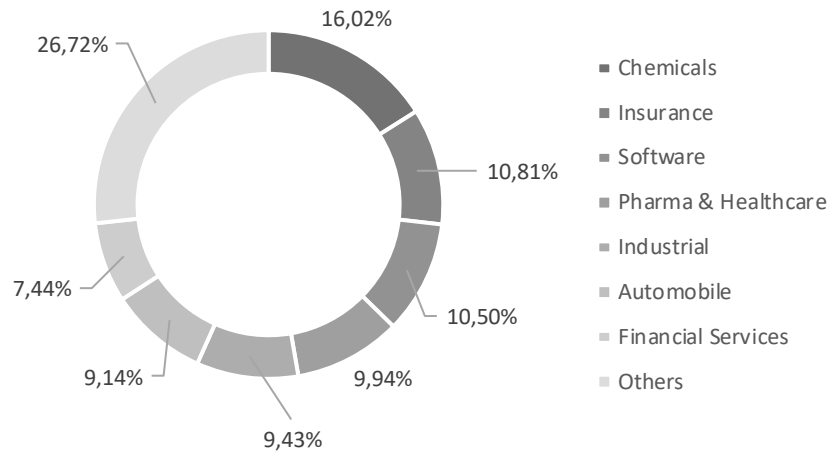


Abbildung 5: Sektorale Zusammensetzung des DAX

Der DAX zählt als Performance Index zu einer der wenigen Indizes, die auch Dividendenerträge abbilden und somit die gesamte aktuelle Performance eines Investments widerspiegeln (Dax-indices.com., 2020c). Da der DAX der bekannteste deutsche Aktienindex ist, wird er häufig als Leitindex sowie repräsentativ für den deutschen Aktienmarkt erachtet und gilt als Stimmungsbarometer der deutschen Wirtschaft (Dax-indices.com., 2020a). Seit seiner Einführung im Jahr 1988 von der Deutschen Börse wird der DAX quartalsweise neu berechnet, neue Unternehmen werden aufgenommen, andere wiederum ausgeschlossen. Als Grundlage für die Neuberechnung werden die Xetra-Kurse des elektronischen Handelsplatzes der Frankfurter Wertpapierbörse herangezogen (Kühn und Kühn, 2019, S. 87).

Zu den Voraussetzungen für eine potentielle Aufnahme in den DAX zählen eine Listung im Prime Standard, ein kontinuierlicher Handel über Xetra sowie ein Streubesitzanteil von zumindest 10 %. Darüberhinausgehend ist das Unternehmen in Deutschland ansässig oder der Großteil des Aktienhandels findet an der Frankfurter Börse statt, und die Aktiengesellschaft sitzt in der EU. Die Auswahl aus jenen Unternehmen, die diese Mindestanforderungen erfüllen, erfolgt entsprechend der Kriterien Marktkapitalisierung und Börsenumsatz. Da diese einer stetigen Veränderung unterliegen, wird jährlich zum Anpassungstermin im September ermittelt, ob die Zusammensetzung des DAX noch den tatsächlichen Marktverhältnissen entspricht, gegebenenfalls werden Anpassungen vorgenommen (Unternehmen aufgenommen oder ausgeschlossen). Im Zuge von

Insolvenzen und Fusionen können auch außertourlich Veränderungen der Zusammensetzung des DAX erforderlich werden. Zusätzlich gilt, dass die Aktien einer Gesellschaft lediglich bis zu 10 % am Index ausmachen dürfen, wobei die Streubesitz-Marktkapitalisierung und nicht der gesamte Börsenwert entscheidend ist (Kühn und Kühn, 2019, S. 87).

Bis September 2021 soll eine DAX-Reform umgesetzt werden, die Aufnahme zusätzlicher Qualifikationskriterien und eine Harmonisierung mit internationalen Standards vorsieht. So ist unter anderem eine Listung am Prime Standard künftig nicht mehr zwingend erforderlich, stattdessen wird es eine Verpflichtung zur Publikation jährlicher Finanzberichte sowie Quartalsberichte geben. Im Rahmen der Reform soll auch die Zusammensetzung des DAX auf 40 Unternehmen ansteigen (Xetra.com., 2020).

Die untenstehende Abbildung 6 zeigt die Entwicklung des DAX im Betrachtungszeitraum 2000 bis 2019 (Finance.yahoo.com., 2020).



Abbildung 6: Entwicklung des DAX 2000-2019

In diesem Zeitraum stieg der DAX mit einer durchschnittlichen Jahresrendite von 5,5 % von 7.396 Punkten auf 13.432 Punkte. Der Index weist somit eine solide jährliche Rendite auf, die verglichen mit anderen Indizes zudem eine geringe Volatilität zeigt. Des Weiteren ist der DAX einer der weltweit meistgehandelten Aktienindexderivate und daher sehr

liquide. Ein zusätzlicher Vorteil des DAX ist dessen Diversifikationsgrad. Neben der Automobilbranche sind auch die Chemie-, Banken- und Industriesparte im DAX vertreten. Die Regelung, die die Marktkapitalisierung eines Unternehmens auf maximal 10 % („Capping Factor“) am Gesamtindex beschränkt, stellt die Diversifikation des DAX sicher, da der Index nicht durch eine einzelne Gesellschaft dominiert wird (Dax-indices.com., 2020c).

5.3.2 S&P 500®

Der Standard & Poor's 500 (in weiterer Folge als „S&P 500“ abgekürzt) umfasst die 505 größten US-amerikanischen Aktiengesellschaften, die an der New Yorker Börse gelistet sind: Der S&P 500 ist neben S&P Midcap 400 und S&P Smallcap 600 Bestandteil des S&P Composite 1500. Seine Ursprünge gehen auf das Jahr 1923 zurück, als das damalige Standard Statistics wöchentlich einen Index publizierte, der 233 Unternehmen aus 26 verschiedenen Industrien umfasste (spglobal.com., 2021). Der S&P 500, so wie dieser heute bekannt ist, wurde 1957 eingeführt und war damals der erste nach Marktkapitalisierung gewichtete Aktienindex in US-Amerika. Diese Indexgewichtung hat zur Folge, dass der S&P 500 in der klassischen Variante als Kursindex einzig durch die Höhe der Aktienkurse beeinflusst wird und nicht durch Dividendenerträge oder Zinszahlungen. Es ist so zwar keine langfristige Aussage über dessen Performance, wohl aber über die Kursentwicklung der Aktien möglich (Zantow und Dinauer, 2011, S. 239).

Für die Aufnahme eines Unternehmens in den S&P 500 ist die Profitabilität das Hauptkriterium (Kühn und Kühn, 2019, S. 90). Zu den weiteren Voraussetzungen, die ein Unternehmen für ein Listing im S&P 500 erfüllen muss, zählen unter anderem eine hohe Liquidität der Aktien, ein Streubesitzanteil von zumindest 10 % und eine Minimum-Marktkapitalisierung von USD 9,8 Milliarden. Des Weiteren muss es sich um ein US-amerikanisches Unternehmen handeln, das im letzten Quartal einen Gewinn erwirtschaftet hat, sowie ein positives Ergebnis über den Zeitraum der letzten vier vorangegangenen Quartale aufweisen kann. Bei der Listung im S&P 500 wird zudem darauf geachtet, dass das Verhältnis zwischen den verschiedenen Sektoren – gemessen an der Marktkapitalisierung – relativ ausgewogen ist (spglobal.com., 2021).

Die nachfolgende Tabelle 12 umfasst einen Auszug der im S&P 500 gelisteten Unternehmen. Dargestellt werden die (gemessen an ihrer Marktkapitalisierung) 10 größten Unternehmen mit Stand 31.12.2020 (spglobal.com., 2021):

Instrument	Symbol
APPLE INC.	APPL
MICROSOFT CORP	MSFT
AMAZON.COM INC	AMZN
FACEBOOK INC A	FB
TESLA, INC	TSLA
APLPHABET INC A	GOOGL
APLPHABET INC C	GOOG
BERKSHIRE HATHAWAY B	BRK.B
JOHNSON & JOHNSON	JNJ
JP MORGAN CHASE & CO	JPM

Stand 31.12.2020

Tabelle 12: Top 10 der größten S&P 500 Unternehmen

Im S&P 500 sind die 11 GICS ® Sektoren („Global Industry Classification Standards“) Energy, Materials, Industrials, Consumer Discretionary, Consumer Staples, Health Care, Financials, Information Technology, Communication Services, Real Estate und Utilities vertreten. Die nachfolgende Übersicht zeigt die Repräsentation dieser Sektoren im Index. Wie aus der Grafik in Abbildung 7 ersichtlich ist, ist mit einem Anteil von 27,6 % der Großteil der gelisteten Unternehmen der Branche „Information Technology“ zuzuordnen, wobei wiederum rund 4,0 % der Indexgewichtung auf das Unternehmen Apple Inc. entfallen. Mit 13,5 % bzw. 12,7 % sind auch die Sektoren „Health Care“ und „Consumer Discretionary“ stark im Index vertreten (spglobal.com., 2021).

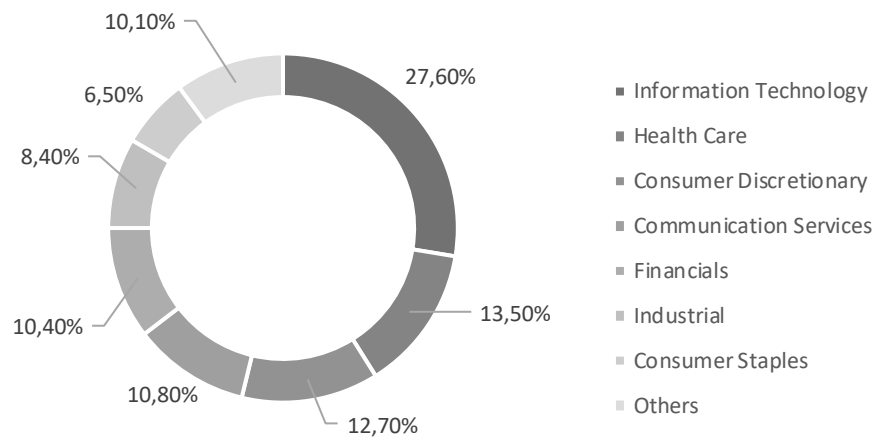


Abbildung 7: Sektorale Zusammensetzung des S&P 500

Da es sich beim S&P 500 um einen entsprechend der Marktkapitalisierung gewichteten Index handelt, unterliegen die Sektor-Anteile aufgrund der veränderlichen Anzahl an sich im Umlauf befindlichen Aktien sowie Kursschwankungen allerdings einer stetigen Veränderung (spglobal.com., 2021).

Die untenstehende Abbildung 8 zeigt die Entwicklung des S&P 500 im Betrachtungszeitraum 2000 bis 2019. In der dargestellten Periode stieg der US-amerikanische Index von 1.402 auf 3.289 Indexpunkte, wobei der Index vor allem seit 2009 ein kontinuierlich starkes Wachstum erlebt. Dies entspricht einem jährlichen Zuwachs von 5,7 % (Finance.yahoo.com., 2021).

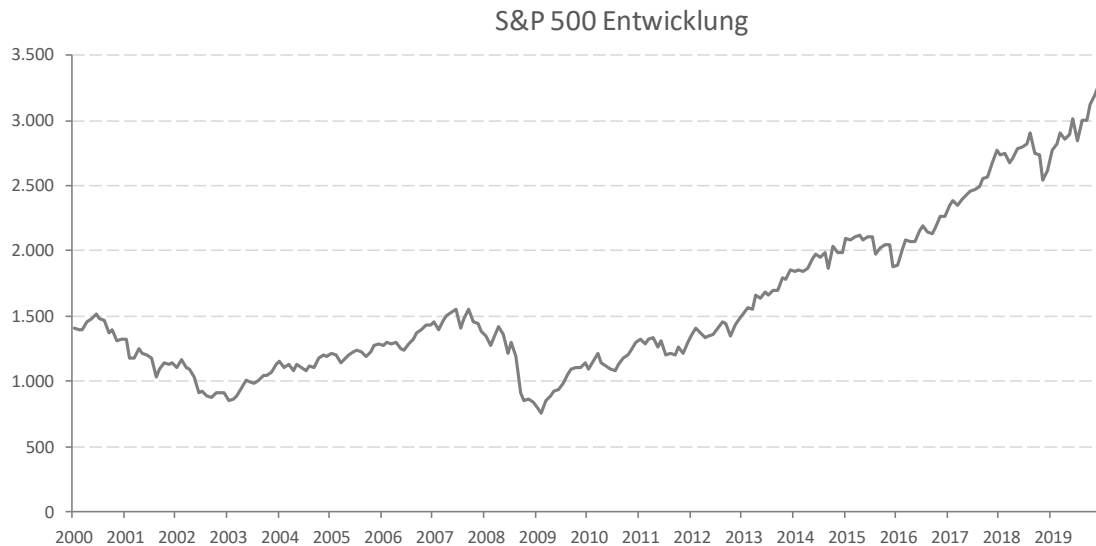


Abbildung 8: Entwicklung des S&P 500 von 2000-2019

6 Erläuterung und formale Beschreibung der Handelsstrategien

Auf Basis der a-posteriori Wahrscheinlichkeiten für einen Anstieg beziehungsweise Rückgang des betrachteten Aktienindex werden Handelssignale abgeleitet und Investitionsentscheidungen getroffen, mit dem Ziel, die Performance des jeweiligen Aktienindex zu übertreffen. Die a-posteriori Wahrscheinlichkeiten wurden mithilfe des Naive Bayes Klassifikators für die relevante Periode generiert. Insgesamt werden vier unterschiedliche Strategien angewendet, die nachfolgend im Detail erläutert werden.

Diesen Strategien ist unter anderem gemein, dass sie allesamt auf den abgeleiteten a-posteriori Wahrscheinlichkeiten beruhen. Damit die Analysen nicht durch den Monatsanfangs- beziehungsweise Monatsendeffekt („Turn-of-the-month Effekt“) verzerrt werden, werden die Investitionsentscheidungen am 15. eines Monats für den Folgemonat gefällt. Arsad et al. (2011, S. 1) definieren den Turn-of-the-month Effekt als: „a situation where stocks perform well and reach a new high level during a period of time that usually consists of the last few days of a month and the first few days of the next month.“ (Arsad et al., 2011, S. 1).

Laut Ziemba (2012, S. 58) betrifft der Turn-of-the-month Effekt im historischen Zeitverlauf die Renditen auf Aktienmärkten zum Monatswechsel (-1 Tag bis +4 Tage), welche in

diesem Zeitraum tendentiell höher ausfallen. Empirische Analysen haben gezeigt, dass der Großteil des monatlichen Wertzuwachses zum Monatswechsel beziehungsweise in den ersten zwei Wochen stattfindet.

Auch wenn die Gründe für den beschriebenen Effekt vielfältig sind, so ist er im Wesentlichen durch das Investitionsverhalten von Institutionen und allgemeinen Zahlungsströmen getrieben. In den USA ist es beispielsweise üblich, dass Zahlungen betreffend Löhne, Gehälter sowie die Begleichung von Schulden einen Tag vor dem Monatsende (-1 Tag) fließen. Da den Menschen ab diesem Zeitpunkt wieder vermehrt liquide Mittel zur Verfügung stehen, ist die Nachfrage am Aktienmarkt bedingt durch das Kaufverhalten von Privatpersonen aber auch institutionellen Anlegern wie Pensionsfonds besonders hoch. Die verstärkte Nachfrage am Aktienmarkt wiederum führt dazu, dass die Renditen zum Monatsanfang beziehungsweise Monatsende über dem Durchschnitt liegen. Abhängig vom jeweiligen Monat variieren die Zahlungsströme und somit auch die Auswirkung auf den Aktienmarkt. Besonders stark zeigt sich der Turn-of-the-month-Effekt im Jänner, da in diesem Monat die höchsten Zahlungsströme fließen. Der Turn-of-the-month-Effekt wird zudem durch verhaltensökonomische Aspekte verstärkt, da positive Nachrichten meist zu Monatsbeginn publiziert werden, schlechte Nachrichten hingegen werden eher später bekanntgegeben (Ziembra, 2012, S. 58f).

Der Investitionszeitraum beträgt jeweils ein Monat bevor erneut eine Entscheidung getroffen wird. Der Investitionsbetrag beim erstmaligen Investment beträgt 100 Geldeinheiten und entwickelt sich entsprechend der erzielten Renditen – der Investitionsbetrag ist somit veränderlich. Erzielte Gewinne werden demnach weder realisiert noch werden Verluste ausgeglichen. Dies gilt auch für den Fall, dass das anfangs investierte Kapital durch Verluste zur Gänze aufgebraucht werden sollte.

6.1 Strategie 1 (Simple Strategie)

Bei Anwendung von Strategie 1 wird für den Fall, dass die a-posteriori Wahrscheinlichkeit für einen Anstieg des Aktienindex im relevanten Investitionszeitraum mit $> 50\%$ erwartet wird, mit der vollen Höhe des zum Investitionszeitpunkt (= Entscheidungszeitpunkt) t vorhandenen Budgets in den Index investiert (Long Position im Markt). Diese Position wird so lange gehalten und somit eine Rendite entsprechend der tatsächlichen

Entwicklung des Index erzielt, bis die a-posteriori Wahrscheinlichkeit für einen weiteren Anstieg des Index auf $\leq 50\%$ fällt. Sobald dies der Fall ist, werden bestehende Long-Positionierungen im Markt verkauft und die Strategie befindet sich im Geld. Während sich die Strategie im Geld befindet, wird folglich eine Rendite von 0% erwirtschaftet. Anschließend wird die Strategie 1 formal beschrieben:

Schritt 1: Formulierung der Handelsentscheidung

Im Schritt 1 wird auf Basis der a-posteriori Wahrscheinlichkeiten zum Zeitpunkt t die Handelsentscheidung getroffen. Diese wird in Analogie zu Abbildung 9 abgeleitet:

$P[DAX_{t+1} > DAX_t]$ (a-posteriori Wslk Aktienindex steigt (t)) $> 50 \%$			
ja			nein
aktuell Long im Markt		aktuell Long im Markt	
ja	nein	ja	nein
keine Handlung, Long Position beibehalten	Long Instrument kaufen	Long Instrument verkaufen, Positionierung im Geld einnehmen	keine Handlung, im Geld bleiben

Abbildung 9: Ableitung der Handelsentscheidung für Strategie 1

Die Ableitung der Handelsentscheidung wird für jeden Entscheidungszeitpunkt durchgeführt und anschließend anhand eines Zahlenbeispiels in Tabelle 13 veranschaulicht.

(1) Zeitpunkt der Investition t	(2) $P[DAX_{t+1} > DAX_t]$	(3) $P[DAX_{t+1} < DAX_t]$	(4) Positionierung vor Handelsentscheidung	(5) Handelsentscheidung	(6) Positionierung nach Handelsentscheidung
1	20,00%	80,00%	im Geld*	Keine Handlung	im Geld
2	55,00%	45,00%	im Geld	Long Instrument kaufen	Long im Markt
3	60,00%	40,00%	Long im Markt	Keine Handlung	Long im Markt
4	30,00%	70,00%	Long im Markt	Long Instrument verkaufen	im Geld
5	80,00%	20,00%	im Geld	Long Instrument kaufen	Long im Markt
6	20,00%	80,00%	Long im Markt	Long Instrument verkaufen	im Geld

Tabelle 13: Ableitung der Handelsentscheidung für Strategie 1 - Beispiel

*zum Zeitpunkt t = 0 war keine Position im Markt aktiv und das gesamte Portfolio war im Geld

Schritt 2: Ermittlung der erzielten Rendite

Anhand der in Tabelle 13 abgeleiteten Handelsentscheidung und Positionierung wird die Rendite für die aktuelle Periode t nach der Vorgehensweise in Abbildung 10 ermittelt.

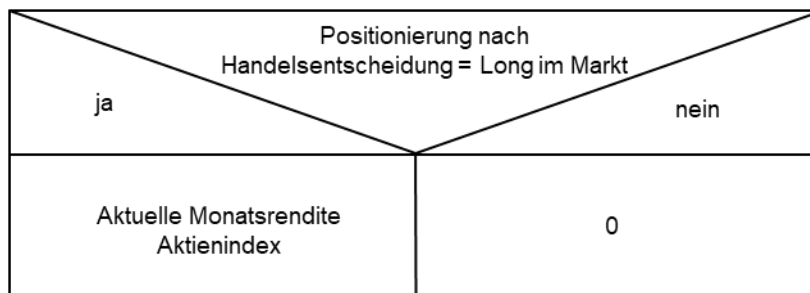


Abbildung 10: Ermittlung der Rendite für Strategie 1

Für eine illustrative Darstellung der erläuterten Strategie 1 für das Jahr 2017 sowie die erzielte Rendite sei an dieser Stelle auf Unterabschnitt 7.3.1 verwiesen.

6.2 Strategie 2 (Durchschnitt für Entscheidung („DFE“), exkl. Short Positionierung)

Im Gegensatz zu Strategie 1, die lediglich das Signal der aktuellen Periode in die Investitionsentscheidung einfließen lässt, werden bei Strategie 2 für die Entscheidungsfindung im Zeitpunkt t alle von Beginn der Untersuchungsperiode bis zum Monat t-1 erhaltenen a-posteriori Wahrscheinlichkeiten > 50 % über einen laufenden Durchschnitt in die Handelsentscheidung einbezogen.

Konkret werden mit den im Entscheidungszeitpunkt t verfügbaren a-posteriori Wahrscheinlichkeiten für einen steigenden beziehungsweise fallenden Aktienkurs, sofern diese $> 50 \%$ waren, fortlaufende Mittelwerte gebildet (= Durchschnitt für Entscheidung; „DFE“ | Anm.: eigene Bezeichnung des Autors). Nur in den Fällen, in denen die a-posteriori Wahrscheinlichkeit für einen steigenden Kurs zum Zeitpunkt t den fortlaufenden Mittelwert der historischen Wahrscheinlichkeiten übersteigt, wird ein Kaufsignal generiert. Selbes gilt auch für Verkaufssignale: Übersteigt in der relevanten Periode die a-posteriori Wahrscheinlichkeit für einen fallenden Index den Mittelwert der historischen a-posteriori Wahrscheinlichkeiten, so wird ein Verkaufssignal abgeleitet. Für den Fall, dass die a-posteriori Wahrscheinlichkeit der relevanten Periode t weder den Mittelwert der historischen a-posteriori Wahrscheinlichkeiten für einen steigenden noch für einen fallenden Kurs übersteigt, wird die Handelsentscheidung der Vorperiode beibehalten.

Wie auch in Strategie 1 werden ebenso bei Strategie 2 bei Verkaufssignalen bestehende Long-Positionierungen im Markt verkauft und die Strategie befindet sich im Geld. Während sich die Strategie im Geld befindet, wird folglich eine Rendite von 0% erwirtschaftet.

Nachfolgend wird die Strategie 2 formal erläutert:

Schritt 1: Ermittlung der laufenden Durchschnitte

$$\emptyset P [DAX_t > DAX_{t-1}] = \sum_{\tau=0}^{t-1} \frac{1}{ANZAHL^1} * \begin{cases} P [DAX_{\tau} > DAX_{\tau-1}]^2, & \text{falls } P [DAX_{\tau} > DAX_{\tau-1}] > 0,5 \\ n. a., & \text{sonst} \end{cases} \quad (36)$$

$$\emptyset P [DAX_t < DAX_{t-1}] = \sum_{\tau=0}^{t-1} \frac{1}{ANZAHL} * \begin{cases} P [DAX_{\tau} < DAX_{\tau-1}], & \text{falls } P [DAX_{\tau} < DAX_{\tau-1}] > 0,5 \\ n. a., & \text{sonst} \end{cases} \quad (37)$$

¹ „ANZAHL“ zählt wie oft $P [DAX_{\tau} > DAX_{\tau-1}] > 0,5$ zutrifft. Das bedeutet beispielsweise, wenn die a-posteriori Wahrscheinlichkeit drei Mal größer 50% ist und zwei Mal kleiner als 50% ist ergibt „ANZAHL“ den Wert drei.

² $P [DAX_{\tau} > DAX_{\tau-1}]$ bezeichnet die a-posteriori Wahrscheinlichkeit eines steigenden Aktienindex zum Zeitpunkt τ . | Für $\tau = 0$ wird jeweils eine a-posteriori Wahrscheinlichkeit von 50% für einen fallenden beziehungsweise steigenden Aktienindex angenommen.

Die formale Darstellung in Gleichung (36) und (37) wird in folgender Tabelle 14 anhand eines Zahlenbeispiels für den Zeitpunkt $t = 6$ dargestellt. Es ergibt sich somit eine durchschnittliche Wahrscheinlichkeit für „Aktienindex steigt“ zum Zeitpunkt $t = 6$ in Höhe von 65,00 % in Spalte 6 $((55,00 \% + 60,00 \% + 80,00 \%) / 3)$. Die durchschnittliche Wahrscheinlichkeit für „Aktienindex fällt“ zum Zeitpunkt $t = 6$ ergibt sich analog in Höhe von 75,00 % in Spalte 7 aus $(80,00 \% + 70,00 \%) / 2$.

(1) Zeitpunkt der Investition t	(2) $P[DAX_{t+1} > DAX_t]$ (a-posteriori Wslk. Aktienindex steigt)	(3) $P[DAX_{t+1} < DAX_t]$ (a-posteriori Wslk. Aktienindex fällt)	(4) $P[DAX_{t+1} > DAX_t] > 50 \%$	(5) $P[DAX_{t+1} < DAX_t] > 50 \%$	(6) $\emptyset P[DAX_t > DAX_{t-1}]$	(7) $\emptyset P[DAX_t < DAX_{t-1}]$
1	20,00%	80,00%	n.a.	80,00%	50,00%*	50,00%*
2	55,00%	45,00%	55,00%	n.a.	50,00%**	80,00%
3	60,00%	40,00%	60,00%	n.a.	55,00%	80,00%
4	30,00%	70,00%	n.a.	70,00%	57,50%	80,00%
5	80,00%	20,00%	80,00%	n.a.	57,50%	75,00%
6	20,00%	80,00%	n.a.	80,00%	65,00%	75,00%

Tabelle 14: Ermittlung der laufenden Durchschnitte für Strategie 2 - Beispiel

*Zum Zeitpunkt $t = 1$ wird annahmegemäß von durchschnittlichen Wahrscheinlichkeiten für einen Anstieg sowie einen Rückgang des Aktienindex in Höhe von 50 % ausgegangen.

**Da im Zeitpunkt $t = 1$ die a-posteriori Wahrscheinlichkeit für einen steigenden Aktienindex nicht $> 50 \%$ ist und somit kein Durchschnitt errechnet werden kann werden die 50 % der Vorperiode beibehalten.

Schritt 2: Formulierung der Handelsentscheidung

Im nächsten Schritt wird auf Basis der a-posteriori Wahrscheinlichkeiten zum Zeitpunkt t und der durchschnittlichen Wahrscheinlichkeiten zum Zeitpunkt t die Handelsentscheidung getroffen. Diese wird in Analogie zu Abbildung 11 getroffen:

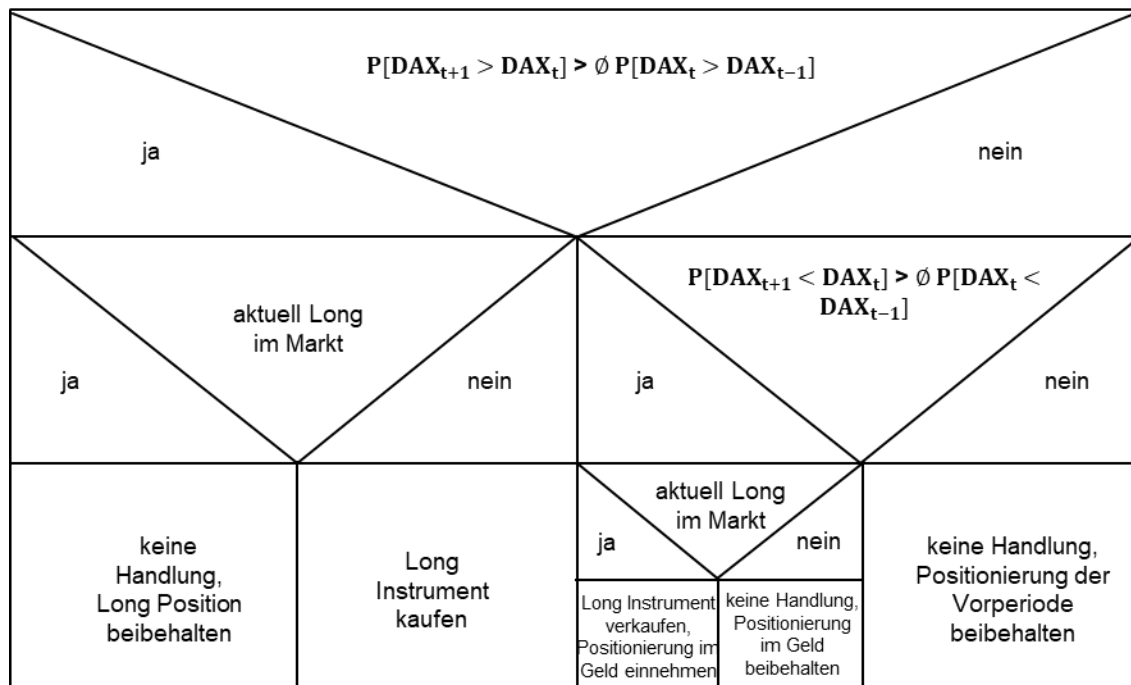


Abbildung 11: Ableitung der Handelsentscheidung für Strategie 2

Die Ableitung der Handelsentscheidung wird für jeden Entscheidungszeitpunkt durchgeführt und anschließend anhand eines Zahlenbeispiels in Tabelle 15 veranschaulicht.

(1) Zeitpunkt der Investition t	(2) $P[DAX_{t+1} > DAX_t]$	(3) $P[DAX_{t+1} < DAX_t]$	(4) $\emptyset P[DAX_t > DAX_{t-1}]$	(5) $\emptyset P[DAX_t < DAX_{t-1}]$	(6) Positionierung vor Handelsentscheidung	(7) Handelsentscheidung	(8) Positionierung nach Handelsentscheidung
1	20,00%	80,00%	50,00%	50,00%	im Geld*	Keine Handlung	im Geld
2	55,00%	45,00%	50,00%	80,00%	im Geld	Long Instrument kaufen	Long im Markt
3	60,00%	40,00%	55,00%	80,00%	Long im Markt	Keine Handlung	Long im Markt
4	30,00%	70,00%	57,50%	80,00%	Long im Markt	Keine Handlung	Long im Markt
5	80,00%	20,00%	57,50%	75,00%	Long im Markt	Keine Handlung	Long im Markt
6	20,00%	80,00%	65,00%	75,00%	Long im Markt	Long Instrument verkaufen	im Geld

Tabelle 15: Ableitung der Handelsentscheidung für Strategie 2 - Beispiel

*zum Zeitpunkt t = 0 wurde keine Position im Markt eingenommen und das gesamte Portfolio war im Geld

Schritt 3: Ermittlung der erzielten Rendite

Anhand der in Tabelle 15 abgeleiteten Handelsentscheidung und Positionierung wird die Rendite für die aktuelle Periode t nach der Vorgehensweise in Abbildung 12 ermittelt.

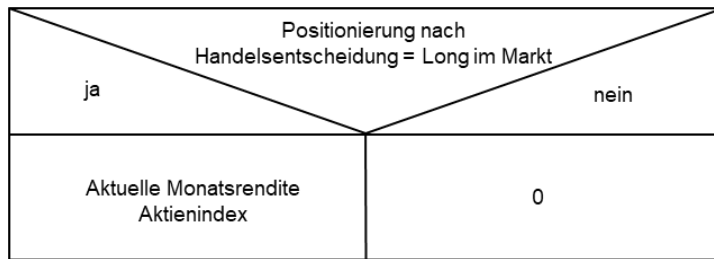


Abbildung 12: Ermittlung der Rendite für Strategie 2

Für eine illustrative Darstellung der erläuterten Strategie 2 für das Jahr 2017 sowie die erzielte Rendite sei an dieser Stelle auf Unterabschnitt 7.3.2 verwiesen.

6.3 Strategie 3 (Durchschnitt für Entscheidung („DFE“), inkl. Short Positionierung)

Wie in Strategien 1 und 2 werden die a-posteriori Wahrscheinlichkeiten für einen Anstieg beziehungsweise einen Rückgang des Aktienindex mithilfe des Naive Bayes Klassifikators für die entsprechende Periode generiert. Die Beurteilung des Umfelds des Aktienmarkts (das bedeutet, ob der Aktienindex steigen oder fallen soll) wird anhand der a-posteriori Wahrscheinlichkeit der aktuellen Periode getroffen.

Diese Strategie bezieht für die Entscheidungsfindung im Zeitpunkt t alle von Beginn der Untersuchungsperiode bis zum Monat t-1 erhaltenen a-posteriori Wahrscheinlichkeiten > 50 % über einen laufenden Durchschnitt analog zu Strategie 2 in die Handelsentscheidung ein.

Analog zu Strategie 2 wird nur in den Fällen, in denen die a-posteriori Wahrscheinlichkeit für einen steigenden Kurs zum Zeitpunkt t den fortlaufenden Mittelwert der historischen Wahrscheinlichkeiten übersteigt ein Kaufsignal generiert (analoge Vorgehensweise für Verkaufssignale wie in Strategie 2).

Wird in der Folge eine Verkaufsentscheidung getroffen, wird eine etwaige bestehende Long-Positionierung im Markt geschlossen und es wird das gesamte Kapital in eine Position investiert, bei welcher von fallenden Aktienindexkursen profitiert wird (Short-Positionierung im Markt). Hierfür wird ein Produkt gekauft, das die inverse Entwicklung des Aktienindex abbildet (beispielsweise ein Short-ETF).

Im Falle einer Long-Positionierung im Markt wird in der betrachteten Monatsperiode folglich die Rendite des Aktienindex (eine positive Rendite, wenn dieser steigt, eine negative, wenn dieser fällt) realisiert, im Falle einer Short-Positionierung im Markt (beispielsweise über einen Short-ETF) wird die aktuelle Monatsrendite des Aktienindex multipliziert mit -1 (eine positive Rendite, wenn dieser fällt, eine negative Rendite, wenn dieser steigt) realisiert.

Diese Strategie kann formal wie folgt dargestellt werden:

Schritt 1: Ermittlung der laufenden Durchschnitte

Die Ermittlung der laufenden Durchschnitte wird analog zu Strategie 2 durchgeführt. Für Details sei deshalb auf Tabelle 14 verwiesen.

Schritt 2: Formulierung der Handelsentscheidung

Im nächsten Schritt wird auf Basis der a-posteriori Wahrscheinlichkeiten zum Zeitpunkt t und der durchschnittlichen Wahrscheinlichkeiten zum Zeitpunkt t die Handelsentscheidung getroffen. Diese wird in Analogie zu Abbildung 13 getroffen:

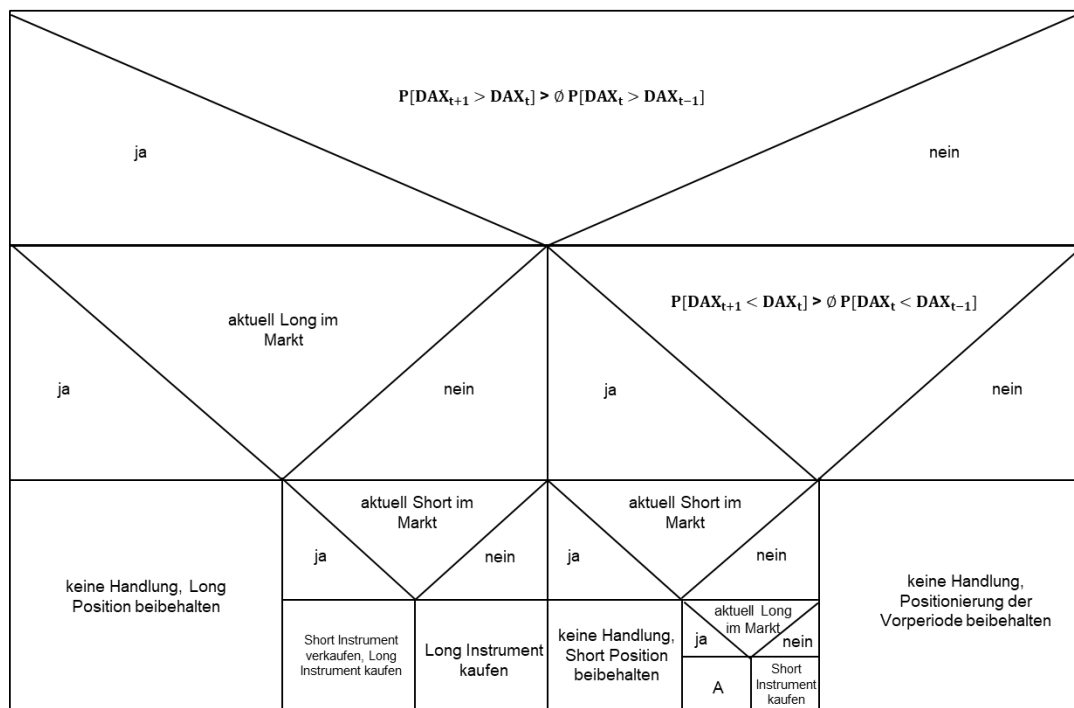


Abbildung 13: Ableitung der Handelsentscheidung für Strategie 3

A = Long Instrument verkaufen, Short Instrument kaufen

Die Ableitung der Handelsentscheidung wird für jeden Entscheidungszeitpunkt durchgeführt und anschließend anhand eines Zahlenbeispiels in Tabelle 16 veranschaulicht.

(1) Zeitpunkt der Investition t	(2) $P[DAX_{t+1} > DAX_t]$	(3) $P[DAX_{t+1} < DAX_t]$	(4) $\emptyset P[DAX_t > DAX_{t-1}]$	(5) $\emptyset P[DAX_t < DAX_{t-1}]$	(6) Positionierung vor Handelsentscheidung	(7) Handelsentscheidung	(8) Positionierung nach Handelsentscheidung
1	20,00%	80,00%	50,00%	50,00%	im Geld*	Short Instrument kaufen	Short im Markt
2	55,00%	45,00%	50,00%	80,00%	Short im Markt	Short Instrument verkaufen, Long Instrument kaufen	Long im Markt
3	60,00%	40,00%	55,00%	80,00%	Long im Markt	Keine Handlung	Long im Markt
4	30,00%	70,00%	57,50%	80,00%	Long im Markt	Keine Handlung	Long im Markt
5	80,00%	20,00%	57,50%	75,00%	Long im Markt	Keine Handlung	Long im Markt
6	20,00%	80,00%	65,00%	75,00%	Long im Markt	Long Instrument verkaufen, Short Instrument kaufen	Short im Markt

Tabelle 16: Ableitung der Handelsentscheidung für Strategie 3 - Beispiel

*zum Zeitpunkt $t = 0$ wurde keine Position im Markt eingenommen und das gesamte Portfolio war im Geld

Schritt 3: Ermittlung der erzielten Rendite

Anhand der in Tabelle 16 abgeleiteten Handelsentscheidung und Positionierung wird die Rendite für die aktuelle Periode t nach der Vorgehensweise in Abbildung 14 ermittelt.

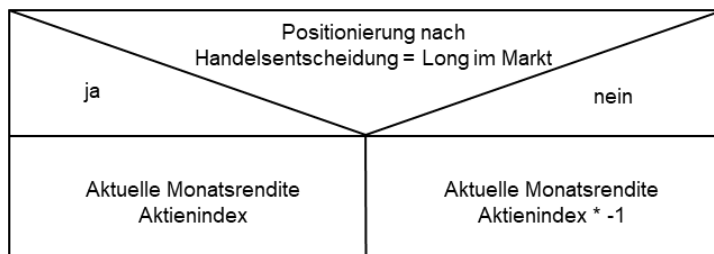


Abbildung 14: Ermittlung der Rendite für Strategie 3

Eine illustrative Darstellung der Strategie 3 für das Jahr 2017 wird in Unterabschnitt 7.3.3 gegeben.

6.4 Strategie 4 (Historisch optimale Entscheidung („HOE“), keine Short Positionierung)

In Strategie 4 wird anders als in Strategie 2 und 3 nicht auf die durchschnittliche historische a-posteriori Wahrscheinlichkeit von $t = 0$ bis zur jeweiligen Vorperiode $(t-1)$ abgestellt. Stattdessen basiert Strategie 4 auf Informationen, die aus dem jeweiligen vorherigen Kalenderjahr gewonnen werden. Mithilfe einer Datentabelle werden für das

jeweils vorangegangene Kalenderjahr ex-post jene Grenzwerte für die a-posteriori Wahrscheinlichkeiten für einen ansteigenden Aktienindex und einen fallenden Aktienindex gesucht, mithilfe derer die Rendite im jeweiligen Jahr maximiert hätte werden können.

Diese Grenzwerte werden in 5 %-Schritten analysiert und zu jeder Kombination von Grenzwerten wird die Rendite ermittelt, die sich auf Basis dieser Grenzwerte für die a-posteriori Wahrscheinlichkeiten ergeben hätte. Die auf diese Weise gewonnen Informationen werden in der aktuell betrachteten Periode in die Generierung der Handelsentscheidung einbezogen. Ergibt die Analyse des Vorjahres beispielsweise, dass bei einer Kaufentscheidung basierend auf einer a-posteriori Wahrscheinlichkeit für den Anstieg des Aktienindex $> 70\%$ („kritisches Niveau Kaufsignal Vorperiode (t-1)“) die beste Performance erzielt werden konnte, so wird in der aktuellen Periode nur ein Kaufsignal generiert, sofern die aktuelle a-posteriori Wahrscheinlichkeit für den Anstieg des Aktienindex bei $> 70\%$ liegt. Das kritische Niveau gilt jeweils für ein Kalenderjahr. Die Methodik wird sinngemäß für ein Verkaufssignal angewandt, wobei innerhalb dieser Strategie entweder eine Long-Positionierung im Markt eingegangen werden kann oder es befindet sich das Portfolio im Geld. Für den Fall, dass zur Zeit einer bestehenden Long-Positionierung im Markt ein Verkaufssignal generiert wird, wird die bestehende Long-Positionierung im Markt verkauft und Cash gehalten (das Portfolio befindet sich im Geld), solange bis erneut ein Kaufsignal generiert wird.

Die vorgestellte Strategie lässt sich formal wie folgt beschreiben:

Schritt 1: Ermittlung der kritischen Werte

Die Ermittlung der kritischen Werte erfolgt wie oben erläutert mithilfe einer Datentabelle. Diese Ermittlung ist in Unterabschnitt 7.3.4 beispielhaft dargestellt.

Schritt 2: Formulierung der Handelsentscheidung

Im ersten Schritt wird auf Basis der a-posteriori Wahrscheinlichkeiten zum Zeitpunkt t und der kritischen Niveaus für Kauf- und Verkaufssignale der Vorperiode die Handelsentscheidung getroffen. Diese wird in Analogie zu Abbildung 15 getroffen:

Schritt 3: Ermittlung der erzielten Rendite

Anhand der in Tabelle 17 abgeleiteten Handelsentscheidung und Positionierung wird die Rendite für die aktuelle Periode t nach der Vorgehensweise in Abbildung 16 ermittelt.

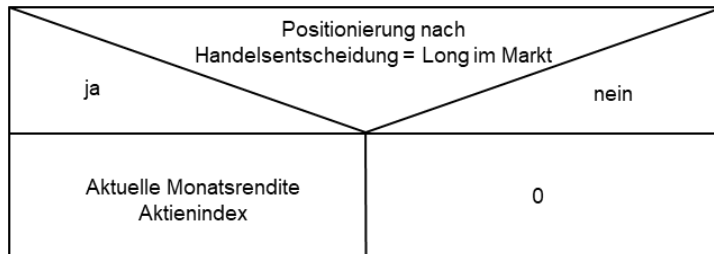


Abbildung 16: Ermittlung der Rendite für Strategie 4

Eine illustrative Darstellung der Strategie 4 für das Jahr 2017 wird in Unterabschnitt 7.3.4 gegeben.

7 Empirische Auswertung der Daten

7.1 Überführung einer Long List in eine Short List

Für den Zeitraum 2013 bis 2016 werden die in Kapitel 4 dargelegten Überlegungen auf den DAX angewendet und die in R durchgeführten Analysen Schritt für Schritt erklärt.

Folgende R-Packages wurden für die nachfolgend beschriebenen Analysen verwendet:

- Library (aTSA) | (Qiu, 2015)
- Library (olsrr) | (Hebbali, 2020)
- Library (car) | (Fox and Weisberg, 2019)
- Library (openxlsx) | (Schaberger and Walker, 2020)

Die statistischen Tests wurden für die makroökonomischen Daten über eine Zeitreihe von vier Jahren („Training Sample“) durchgeführt. Der vierjährige Zeitraum wurde gewählt, da aufgrund des zentralen Grenzwertsatzes gilt, dass sich Stichproben, die einen Umfang > 30 aufweisen einer Normalverteilung annähern (Weigand, 2019, S. 249). Bei monatlichen Messungen über einen vierjährigen Zeitraum umfasst die Stichprobengröße je volkswirtschaftlichen Indikator 48 Datenpunkte, wodurch der zentrale Grenzwertsatz als erfüllt anzusehen ist. Da auch die dem Naive Bayes Klassifikator zugrunde gelegte

Verteilung eine Normalverteilung ist, wird eine konsistente Vorgehensweise sichergestellt.

In einem ersten Schritt wurden die volkswirtschaftlichen Indikatoren mittels ADF-Tests auf Stationarität getestet. Die Lag-Länge wurde nicht einheitlich für sämtliche Indikatoren festgelegt, sondern wurde mit Hilfe einer Daumenregel, welche die Lag-Länge in Abhängigkeit von der Anzahl der Datenpunkte (diese weichen in einzelnen Jahren von den zuvor erwähnten 48 Datenpunkten ab – etwa, wenn in einem oder mehreren Monaten keine Daten publiziert wurden) ermittelt. Dies gilt sowohl für die Lag-Länge der relevanten makroökonomischen Indikatoren sowie für die erklärte Variable, den DAX, im Zeitraum 2013 bis 2016. Um möglichst konsistente Entscheidungen zu treffen, wurden einheitlich für alle Indikatoren die 3 Varianten des ADF-Tests "no drift no trend", "with drift no trend" und "with drift and trend" durchgeführt. Wurde das Signifikanzniveau, das mit einer Höhe von 10 % festgelegt wurde, bei einer der Varianten überschritten, wurde eine (weitere) Differenz des Indikators gebildet und abermals getestet (Qiu, 2015).

In der nachfolgenden Abbildung 17 wird beispielhaft der ADF-Test für die erste Differenz des Indikators „Wechselkurs EUR-USD“ dargestellt. Es ist ersichtlich, dass in Variante 3 „with drift and trend“ das Signifikanzkriterium ($\alpha = 0,1$) verfehlt wird, weshalb für diesen Indikator die zweite Differenz in weiterer Folge herangezogen wurde (Signifikanzkriterium erfüllt).

```
> adf.test(y_EURUSD.diff1[!is.na(y_EURUSD.diff1)])
Augmented Dickey-Fuller Test
alternative: stationary
```

```
Type 1: no drift no trend
      lag   ADF p.value
[1,]    0 -7.10  0.0100
[2,]    1 -3.82  0.0100
[3,]    2 -3.00  0.0100
[4,]    3 -2.21  0.0286
Type 2: with drift no trend
      lag   ADF p.value
[1,]    0 -7.43  0.0100
[2,]    1 -4.00  0.0100
[3,]    2 -3.14  0.0335
[4,]    3 -2.44  0.1632
Type 3: with drift and trend
      lag   ADF p.value
[1,]    0 -7.35  0.0100
[2,]    1 -4.00  0.0178
[3,]    2 -3.16  0.1088
[4,]    3 -2.43  0.3935
----
```

Abbildung 17: Augmented Dickey-Fuller Test in R

Mit jenen Indikatoren beziehungsweise deren Differenzen, die das Stationaritätskriterium erfüllen, wurde eine multiple lineare Regressionsgleichung folgender Form aufgestellt (vgl. Abbildung 18):

```
lm(y_DAX.diff2 ~ y_cpi.diff1+y_zew.curr.sit.diff2+y_zew.exp.diff2+y_ret.sal+y_unemp.ch.diff1
+y_unemp.ra.diff2+y_ifo.clim.diff2+y_ifo.exp.diff2+y_ifo.curr.ass.diff2+y_fact.ord+y_ind.prod
+y_PPI.diff1+y_trad.bal.diff1+y_GfK.diff2+y_Manu.PMI.diff1+y_Serv.PMI.diff1+y_sov.crv.3m.diff2+y_sov.crv.6m.diff2
+y_sov.crv.10y.diff2+y_EURUSD.diff2)
```

Abbildung 18: Multiple lineare Regressionsgleichung in R

Für die Verringerung von Multikollinearität im Modell wurde im nächsten Schritt der VIF-Faktor für die obenstehenden exogenen Variablen ermittelt. Jene Variable, die den höchsten über dem Schwellwert von 10 liegenden Faktor aufweist, wurde zuerst aus dem Modell ausgeschlossen und die VIF-Faktoren erneut kalkuliert. Diese Vorgehensweise wurde so lange wiederholt, bis keine der übrig gebliebenen Regressoren mehr den Schwellwert überschreitet.

Abbildung 19 zeigt den mit R durchgeführten VIF-Test. Die zweite Differenz des Indikators „Ifo Business Climate“ weist den höchsten über 10 liegenden VIF-Wert auf und wird somit als erster Indikator ausgeschlossen.


```
> vif(model) # VIF test
      y_cpi.diff1 y_zew.curr.sit.diff2 y_zew.exp.diff2 y_ret.sal y_unemp.ch.diff1 y_unemp.ra.diff2
      3.441007      3.560000      2.623915      3.570204      2.555156      2.551084
y_ifo.clm.diff2 y_ifo.exp.diff2 y_ifo.curr.ass.diff2 y_fact.ord y_ind.prod y_PPI.diff1
1499.143186      611.664651      446.522893      3.175038      2.200383      3.341933
y_trad.bal.diff1 y_GfK.diff2 y_Manu.PMI.diff1 y_Serv.PMI.diff1 y_sov.crv.3m.diff2 y_sov.crv.6m.diff2
      3.017273      1.733578      2.567959      1.792071      2.220210      2.896113
y_sov.crv.10y.diff2 y_EURUSD.diff2
      3.097679      3.335352
```

Abbildung 19: VIF-Test in R

Auf Basis der um jene Faktoren bereinigten Indikatoren, die das VIF-Kriterium verfehlen, wurde die folgende Regressionsgleichung aufgestellt (vgl. Abbildung 20), die in einem weiteren Schritt als Basis für die Selektionsverfahren „Forward Selection“ und „Backward Elimination“ herangezogen wurde:

```
lm(y_DAX.diff2 ~ y_cpi.diff1+y_zew.curr.sit.diff2+y_zew.exp.diff2+y_ret.sal+y_unemp.ch.diff1
+y_unemp.ra.diff2+y_ifo.exp.diff2+y_ifo.curr.ass.diff2+y_fact.ord+y_ind.prod
+y_PPI.diff1+y_trad.bal.diff1+y_GfK.diff2+y_Manu.PMI.diff1+y_Serv.PMI.diff1+y_sov.crv.3m.diff2+y_sov.crv.6m.diff2
+y_sov.crv.10y.diff2+y_EURUSD.diff2)
```

Abbildung 20: Multiple lineare Regressionsgleichung nach Durchführung VIF-Test

Um jene makroökonomischen Indikatoren auszuwählen, die signifikant für die Erklärung des DAX sind und in weiterer Folge für die Investitionsentscheidungen im „Naive Bayes Investing“ eingesetzt werden, wurden die beiden Selektionsverfahren Forward Selection und Backward Elimination herangezogen, wobei im Rahmen dieser Arbeit für beide Verfahren ein Signifikanzniveau in Höhe von 10 % als angemessen erachtet wurde. Prinzipiell können die beiden Selektionsverfahren zu unterschiedlichen Ergebnissen führen. Um präzisere Ergebnisse zu erhalten, wurden im Rahmen dieser Arbeit nur jene Variablen für das „Naive Bayes Investing“ verwendet, die sowohl mittels Forward Selection als auch Backward Elimination selektiert wurden. Abbildung 21 und Abbildung 22 zeigen die Forward Selection und Backward Elimination in R:

Selection Summary						
Step	Variable Entered	R-Square	Adj. R-Square	C(p)	AIC	RMSE
1	y_ind.prod	0.1534	0.1350	9.0483	764.0366	664.1821
2	y_ifo.curr.ass.diff2	0.3195	0.2893	0.6391	755.5519	602.0435
3	y_zew.curr.sit.diff2	0.3703	0.3273	-0.5389	753.8340	585.7176

Abbildung 21: Ergebnis der Forward Selection Methode

model	Beta	Std. Error	Std. Beta	t	Sig	lower	upper
(Intercept)	1.233	86.942		0.014	0.989	-173.877	176.344
y_ifo.curr.ass.diff2	-160.061	48.292	-0.411	-3.314	0.002	-257.326	-62.796
y_ind.prod	236.411	65.939	0.445	3.585	0.001	103.604	369.219

Abbildung 22: Ergebnis der Backward Elimination Methode

Entsprechend den durchgeführten statistischen Tests und Analysen zeigte sich für den Zeitraum 2013 bis 2016, dass die beiden Indikatoren „Ifo Current Assessment (2. Differenz)“ und „Industrial Production (monatliche Veränderung)“ für die Investitionsentscheidungen im „Naive Bayes Investing“ relevant sind.

7.2 Anwendung des Naive Bayes Klassifikators

Dieser Abschnitt widmet sich nun schlussendlich der Anwendung des Naive Bayes Klassifikators auf die vorgestellten volkswirtschaftlichen Indikatoren und liefert eine Schritt für Schritt Anleitung zur Ermittlung der a-posteriori Wahrscheinlichkeiten, die in weiterer Folge zur Generierung der Handelssignale in der jeweiligen Strategie herangezogen werden. Die zentralen Inputfaktoren stellen die Short Listen für die Jahre 2004 bis 2019 zur Verfügung, deren Ableitung aus der Long List bereits in Kapitel 4 dargestellt wurde. Die in jeder Short List vorkommenden Indikatoren werden analog zur Vorgehensweise in Unterabschnitt 3.2.7 zur Ableitung der a-posteriori Wahrscheinlichkeiten für einen steigenden beziehungsweise fallenden Aktienindex herangezogen. Als „steigender Aktienindex“ wird in dieser Arbeit eine positive Veränderung des Aktienindex gegenüber dem Vormonat (15. des aktuellen Monats gegenüber 15. des Vormonats, sofern dies jeweils Handelstage sind) von $> 0,00\%$ definiert. Als „fallender Aktienindex“ wird eine monatliche Veränderung von $< 0,00\%$ festgelegt.

So ergibt sich für das Jahr 2017 für den deutschen Markt auf Basis der durchgeführten Analysen folgende Short List bestehend aus den zwei volkswirtschaftlichen Indikatoren „Ifo Current Assessment (2. Differenz)“ und „Industrial Production (monatliche Veränderung)“. Diese werden, wie in Abschnitt 7.1 dargestellt, auf Basis zahlreicher statistischer Kriterien als angemessen für die Signalgebung erachtet.

(1) Beobachtung#	(2) Ifo Current Assessment (2. Differenz)	(3) Industrial Production (monatliche Veränderung)	(4) Klasse
Jänner 2013	1,90	0,30	Aktienindex fällt
Februar 2013	1,30	0,00	Aktienindex steigt
März 2013	-2,50	0,50	Aktienindex fällt
April 2013	-2,40	1,20	Aktienindex steigt
Mai 2013	5,50	1,80	Aktienindex fällt
Juni 2013	-3,40	-1,00	Aktienindex steigt
Juli 2013	1,30	2,40	Aktienindex steigt
August 2013	1,20	-1,70	Aktienindex steigt
September 2013	-2,50	1,40	Aktienindex steigt
...	...	...	...
Dezember 2016	0,40	0,40	Aktienindex steigt
Mittelwert (Klasse „Aktienindex steigt“)	-0,16	0,37	
Mittelwert (Klasse „Aktienindex fällt“)	0,30	-0,40	
Standardabweichung (Klasse „Aktienindex steigt“)	1,75	1,31	
Standardabweichung (Klasse „Aktienindex fällt“)	1,95	1,29	

Tabelle 18: Ermittlung Mittelwert und Standardabweichung für Short List

Wie in Tabelle 18 ersichtlich, ergibt sich für den Indikator „Ifo Current Assessment (2. Differenz)“ ein Mittelwert von rund -0,16 und eine Standardabweichung von 1,75 bei steigendem Aktienindex. Bei fallendem Aktienindex ergibt sich ein Mittelwert von rund 0,30 sowie eine Standardabweichung von 1,95. Diese Kennzahlen werden basierend auf 48 Datenpunkten ermittelt, wobei monatliche Daten von Jänner 2013 bis inklusive Dezember 2016 in die Berechnung einfließen. Für den Indikator „Industrial Production (monatliche Veränderung)“ werden Mittelwert und Standardabweichung analog ermittelt. Ausgehend von der Short List und deren Kennzahlen in Tabelle 18, werden in einem weiteren Schritt analog zu Unterabschnitt 3.2.7 die a-posteriori Wahrscheinlichkeiten für die beiden möglichen Klassen für die neuen Beobachtungen des Jahres 2017 ermittelt. Diese Ermittlung ist in Tabelle 19 und Tabelle 20 dargestellt.

(1) Zeitpunkt der Investition t	(2) Ifo Current Assessment (2. Differenz) Ausprägung zum jeweiligen Datum („Ifo“)	(3) Industrial Prod. (monatl. Veränderung) Ausprägung zum jeweiligen Datum („IP“)	(4) $P(Ifo=b_t A= \text{Aktienindex steigt})$ („C“)	(5) $P(IP=b_t A= \text{Aktienindex steigt})$ („D“)	(6) $P(Ifo=b_t A= \text{Aktienindex fällt})$ („E“)	(7) $P(IP=b_t A= \text{Aktienindex fällt})$ („F“)	(8) $C \cdot D \cdot 0,56$ („G“)	(9) $E \cdot F \cdot 0,44$ („H“)
15. Jänner 2017	0,40	0,40	0,21679	0,30457	0,20430	0,25563	0,03714	0,02285
15. Februar 2017	-0,70	-3,00	0,21751	0,01118	0,17937	0,04036	0,00137	0,00317
15. März 2017	1,20	2,80	0,16866	0,05419	0,18391	0,01444	0,00514	0,00116
15. April 2017	-0,60	2,20	0,22103	0,11433	0,18391	0,04096	0,01421	0,00330
15. Mai 2017	0,90	-0,40	0,18991	0,25668	0,19512	0,30903	0,02742	0,02638
15. Juni 2017	0,30	0,80	0,22043	0,28843	0,20457	0,20131	0,03576	0,01802
15. Juli 2017	-1,20	1,20	0,19113	0,24882	0,15219	0,14402	0,02675	0,00959
15. August 2017	0,40	-1,10	0,21679	0,16270	0,20430	0,26625	0,01984	0,02380
15. September 2017	-2,10	0,00	0,12323	0,29295	0,09593	0,29489	0,02031	0,01238
15. Oktober 2017	-0,20	2,60	0,22811	0,07115	0,19796	0,02094	0,00913	0,00181
15. November 2017	2,20	-1,60	0,09181	0,09863	0,12727	0,19993	0,00509	0,01113
15. Dezember 2017	-1,60	-1,40	0,16249	0,12262	0,12727	0,22828	0,01121	0,01271

Tabelle 19: Ermittlung der Likelihoods sowie Multiplikation von Likelihood und Prior (Zählergröße) für 2017

Tabelle 19 folgt einem identen Aufbau wie Tabelle 7. Dabei werden die Likelihoods in den Spalten 4 bis 7 („C“ bis „F“) analog zu Unterabschnitt 3.2.7 ermittelt. Auch die Zählergrößen (Likelihood * Prior) in den Spalten 8 und 9 („G“ und „H“) für die Berechnung der a-posteriori Wahrscheinlichkeiten werden ident zu obigem Kapitel abgeleitet. In Tabelle 20 werden diese in weiterer Folge analog zu obigen Ausführungen normiert und zur Klassifikationsentscheidung herangezogen.

(1) Zeitpunkt der Investition t	(2) G	(3) H	(4) $P[DAX_{t+1} > DAX_t]$	(5) $P[DAX_{t+1} < DAX_t]$	(6) Klassifikationsentscheidung	(7) Positionierung vor Handelsentscheidung	(8) Handelsentscheidung	(9) Positionierung im Aktienindex ab 15. des aktuellen Monats
15. Jänner 2017	0,03714	0,02285	61,91%	38,09%	Aktienindex steigt	im Geld*	Long Instrument kaufen	Long im Markt
15. Februar 2017	0,00137	0,00317	30,16%	69,84%	Aktienindex fällt	Long im Markt	Long Instrument verkaufen	im Geld
15. März 2017	0,00514	0,00116	81,56%	18,44%	Aktienindex steigt	im Geld	Long Instrument kaufen	Long im Markt
15. April 2017	0,01421	0,00330	81,18%	18,82%	Aktienindex steigt	Long im Markt	Keine Handlung	Long im Markt
15. Mai 2017	0,02742	0,02638	50,97%	49,03%	Aktienindex steigt	Long im Markt	Keine Handlung	Long im Markt
15. Juni 2017	0,03576	0,01802	66,50%	33,50%	Aktienindex steigt	Long im Markt	Keine Handlung	Long im Markt
15. Juli 2017	0,02675	0,00959	73,61%	26,39%	Aktienindex steigt	Long im Markt	Keine Handlung	Long im Markt
15. August 2017	0,01984	0,02380	45,47%	54,53%	Aktienindex fällt	Long im Markt	Long Instrument verkaufen	im Geld
15. September 2017	0,02031	0,01238	62,13%	37,87%	Aktienindex steigt	im Geld	Long Instrument kaufen	Long im Markt
15. Oktober 2017	0,00913	0,00181	83,43%	16,57%	Aktienindex steigt	Long im Markt	Keine Handlung	Long im Markt
15. November 2017	0,00509	0,01113	31,39%	68,61%	Aktienindex fällt	Long im Markt	Long Instrument verkaufen	im Geld
15. Dezember 2017	0,01121	0,01271	46,86%	53,14%	Aktienindex fällt	im Geld	Keine Handlung	im Geld

Tabelle 20: Ermittlung der a-posteriori Wahrscheinlichkeiten sowie Positionierung im Index für 2017

*Annahme: am 15. Jänner 2017 befand sich das Portfolio im Geld

Ausgehend von Tabelle 20 lässt sich für die einzelnen Monate des Jahres 2017 erkennen, wann der Aktienindex steigen beziehungsweise fallen soll. Ist das Ergebnis des Klassifikators „Aktienindex steigt“, so wird beispielsweise nach Strategie 1 eine Long-Position im Index aufgebaut beziehungsweise eine bereits bestehende Long-Position gehalten. Ist das Ergebnis hingegen „Aktienindex fällt“, so wird eine bestehende Long-Position abgebaut und das Portfolio befindet sich im Geld. Befindet sich das Portfolio bereits im Geld, ist keine Handlung erforderlich. Auf diese Weise ergibt sich im Jahr 2017 für Strategie 1 die Handelsentscheidung für jeden Monat in der Spalte 8, „Handelsentscheidung“.

Hiervon ausgehend lassen sich nun die erzielten monatlichen Renditen der Strategie 1 des betrachteten Jahres relativ einfach ermitteln. Deren Ableitung ist in der nachfolgenden Tabelle 21 dargestellt:

(1) Zeitpunkt der Investition t	(2) Klassifikationsentscheidung	(3) Positionierung im Aktienindex ab 15. des aktuellen Monats	(4) Erzielte Rendite (15. aktueller Monat bis 15. nachfolgender Monat*)	(5) DAX Rendite (ohne Strategie)	(6) Over-/Underperformance
15. Jänner 2017	Aktienindex steigt	Long im Markt	2,07%	2,07%	0,00%
15. Februar 2017	Aktienindex fällt	im Geld	0,00%	1,83%	-1,83%
15. März 2017	Aktienindex steigt	Long im Markt	-0,08%	-0,08%	0,00%
15. April 2017	Aktienindex steigt	Long im Markt	6,72%	6,72%	0,00%
15. Mai 2017	Aktienindex steigt	Long im Markt	-0,90%	-0,90%	0,00%
15. Juni 2017	Aktienindex steigt	Long im Markt	-0,82%	-0,82%	0,00%
15. Juli 2017	Aktienindex steigt	Long im Markt	-3,26%	-3,26%	0,00%
15. August 2017	Aktienindex fällt	im Geld	0,00%	2,81%	-2,81%
15. September 2017	Aktienindex steigt	Long im Markt	3,87%	3,87%	0,00%
15. Oktober 2017	Aktienindex steigt	Long im Markt	-0,21%	-0,21%	0,00%
15. November 2017	Aktienindex fällt	im Geld	0,00%	0,98%	-0,98%
15. Dezember 2017	Aktienindex fällt	im Geld	0,00%	0,74%	-0,74%

Tabelle 21: Ermittlung der erzielten Renditen der Benchmark sowie mithilfe der Strategie 1

*ist der 15. des Monats kein Handelstag, so wird der nächste Handelstag herangezogen

Die erzielte monatliche Rendite eines jeden Monats ist in Spalte 4 in Tabelle 21 dargestellt. Ist die Klassifikationsentscheidung beispielsweise am 15. Jänner 2017 „Aktienindex steigt“, geht die Strategie Long in den Markt und die Position wird bis zum 15. Februar gehalten (sofern diese Tage jeweils Handelstage sind, siehe auch Anmerkung oben). Daraus ergibt sich für diesen Zeitraum eine Rendite von rund 2,07 %. Da der DAX im selben Zeitraum klarerweise dieselbe Rendite erwirtschaftet, ergibt sich

in diesem Fall keine Out- oder Underperformance des Portfolios basierend auf der Strategie 1 gegenüber der Benchmark DAX. Betrachtet man beispielsweise die Investitionsperiode von 15. Februar 2017 bis 15. März 2017, so ist ersichtlich, dass in dieser Periode basierend auf der Strategie 1 keine Investition getätigt wurde, das Portfolio sich im Geld befindet und die erzielte Rendite folglich 0 % ist (zweite Zeile in Tabelle 21). Im gleichen Zeitraum erwirtschaftete der DAX allerdings eine Rendite von rund 1,83 %, wodurch sich eine Underperformance der Strategie 1 gegenüber der Benchmark von 1,83 % ergibt.

7.3 Anwendung der Handelsstrategien

In diesem Abschnitt wird die tatsächliche Anwendung der in Kapitel 6 formal erläuterten Handelsstrategien anhand des Jahres 2017 beispielhaft für den DAX dargestellt.

7.3.1 Strategie 1 (Simple Strategie)

Die nachfolgend dargestellte Tabelle 22 bildet die Entscheidungsgrundlage für die Handelsstrategie 1, „simple Strategie“. In der Spalte 1 ist der Investitionszeitpunkt angegeben. Wesentliches Entscheidungskriterium für die simple Strategie sind die a-posteriori Wahrscheinlichkeiten für einen steigenden beziehungsweise fallenden Aktienkurs. Diese sind den Spalten 2 und 3 zu entnehmen. Auf Basis der Ausprägungen der Indikatoren, die in Abschnitt 7.1 dargestellt wurden, ergibt sich zum 15.01.2017 unter Anwendung des Naive Bayes Klassifikators (Abschnitt 7.2) beispielsweise eine a-posteriori Wahrscheinlichkeit für einen Anstieg des DAX in der nächsten Periode von rund 62 %. Die Positionierung des Investors im Aktienindex (Spalte 7) ist Resultat der Handelsentscheidung und Spalte 6 zu entnehmen. Entsprechend der simplen Strategie wird in den Index investiert, wenn die Wahrscheinlichkeit für einen steigenden Aktienkurs zum fraglichen Investitionszeitpunkt mehr als 50 % beträgt (Long im Markt). Sollte die Strategie bereits Long im Markt sein, wird die Position gehalten und keine Handlung durchgeführt. Ist die Wahrscheinlichkeit für einen steigenden Aktienindex kleiner als 50 % wird eine bestehende Long-Positionierung im Markt veräußert beziehungsweise eine bestehende Positionierung im Geld beibehalten. Spalte 6 zeigt an, welche Handlungen der Investor gemäß der simplen Strategie vornehmen muss.

Zum 15.01.2017 etwa beträgt die a-posteriori Wahrscheinlichkeit für einen Anstieg des DAX etwa rund 62 % und als notwendige Handlung ergibt sich dem > 50 %-Kriterium entsprechend der Kauf eines Long Instruments. Da im Folgezeitpunkt die a-posteriori Wahrscheinlichkeit für einen sinkenden DAX größer ist, werden die gehaltenen Long-Anteile abgestoßen und die Strategie befindet sich im Geld.

(1) Zeitpunkt der Investition t	(2) $P[DAX_{t+1} > DAX_t]$	(3) $P[DAX_{t+1} < DAX_t]$	(4) Klassifikationsentscheidung	(5) Positionierung vor Handelsentscheidung	(6) Handelsentscheidung	(7) Positionierung im Aktienindex ab 15. des aktuellen Monats
15. Jänner 2017	61,91%	38,09%	Aktienindex steigt	im Geld*	Long Instrument kaufen	Long im Markt
15. Februar 2017	30,16%	69,84%	Aktienindex fällt	Long im Markt	Long Instrument verkaufen	im Geld
15. März 2017	81,56%	18,44%	Aktienindex steigt	im Geld	Long Instrument kaufen	Long im Markt
15. April 2017	81,18%	18,82%	Aktienindex steigt	Long im Markt	Keine Handlung	Long im Markt
15. Mai 2017	50,97%	49,03%	Aktienindex steigt	Long im Markt	Keine Handlung	Long im Markt
15. Juni 2017	66,50%	33,50%	Aktienindex steigt	Long im Markt	Keine Handlung	Long im Markt
15. Juli 2017	73,61%	26,39%	Aktienindex steigt	Long im Markt	Keine Handlung	Long im Markt
15. August 2017	45,47%	54,53%	Aktienindex fällt	Long im Markt	Long Instrument verkaufen	im Geld
15. September 2017	62,13%	37,87%	Aktienindex steigt	im Geld	Long Instrument kaufen	Long im Markt
15. Oktober 2017	83,43%	16,57%	Aktienindex steigt	Long im Markt	Keine Handlung	Long im Markt
15. November 2017	31,39%	68,61%	Aktienindex fällt	Long im Markt	Long Instrument verkaufen	im Geld
15. Dezember 2017	46,86%	53,14%	Aktienindex fällt	im Geld	Keine Handlung	im Geld

Tabelle 22: Ermittlung der a-posteriori Wahrscheinlichkeiten sowie der Positionierung im Index für 2017 nach Strategie 1

*Annahme: am 15. Jänner 2017 befand sich das Portfolio im Geld

Die mit der Strategie 1 beziehungsweise mit der Benchmark erzielten Renditen werden in der anschließend dargestellten Tabelle 23 veranschaulicht.

(1) Zeitpunkt der Investition t	(2) Zeitpunkt der Rendite	(3) Positionierung im Aktienindex ab 15. des aktuellen Monats	(4) Rendite basierend auf Strategie 1	(5) DAX Rendite (ohne Strategie)	(6) Over-/Under-performance
15. Jänner 2017	15. Februar 2017	Long im Markt	2,07%	2,07%	0,00%
15. Februar 2017	15. März 2017	im Geld	0,00%	1,83%	-1,83%
15. März 2017	18. April 2017	Long im Markt	-0,08%	-0,08%	0,00%
15. April 2017	15. Mai 2017	Long im Markt	6,72%	6,72%	0,00%
15. Mai 2017	15. Juni 2017	Long im Markt	-0,90%	-0,90%	0,00%
15. Juni 2017	17. Juli 2017	Long im Markt	-0,82%	-0,82%	0,00%
15. Juli 2017	15. August 2017	Long im Markt	-3,26%	-3,26%	0,00%
15. August 2017	15. September 2017	im Geld	0,00%	2,81%	-2,81%
15. September 2017	16. Oktober 2017	Long im Markt	3,87%	3,87%	0,00%
15. Oktober 2017	15. November 2017	Long im Markt	-0,21%	-0,21%	0,00%
15. November 2017	15. Dezember 2017	im Geld	0,00%	0,98%	-0,98%
15. Dezember 2017	15. Jänner 2018	im Geld	0,00%	0,74%	-0,74%
Jahresrendite			7,27%	14,24%	

Tabelle 23: Darstellung der erzielten Renditen der Benchmark und nach Strategie 1 für 2017

In den Spalten 4 und 5 werden die monatlichen Renditen basierend auf Strategie 1 respektive die Renditen, die sich aus einer einfachen Buy-and-Hold Strategie im DAX ergeben dargestellt. Da im ersten Zeitpunkt der Investition (15.01.2017) die a-posteriori Wahrscheinlichkeit für einen steigenden Index überwiegt, wird ein Long Instrument gekauft und eine Long-Position im Markt eröffnet. Die monatliche Rendite beträgt somit 2,07 % und entspricht jener der Benchmark. Für diesen Investitionszeitpunkt ergibt sich somit nach Strategie 1 weder eine Out- noch Underperformance gegenüber der Benchmark DAX. Im Folgezeitpunkt ist die a-posteriori Wahrscheinlichkeit für einen fallenden DAX > 50 %. Somit wird die bestehende Long-Position im DAX geschlossen und das Portfolio befindet sich im Geld. Die erzielte Rendite nach Strategie 1 beträgt somit 0 %, wohingegen der DAX im selben Zeitraum eine Rendite von 1,83 % verzeichnen konnte. Die Underperformance der Strategie 1 gegenüber dem DAX beträgt folglich 1,83 %. Insgesamt wird mit Strategie 1 im Investitionsjahr 2017 eine Rendite in Höhe von 7,27 % erzielt, die deutlich unter jener Rendite liegt, die mit einer einfachen Buy-and-Hold Strategie (14,24 %) erzielt worden wäre. Bei der Ermittlung der Rendite wird der Zinseszins Effekt berücksichtigt.

7.3.2 Strategie 2 (DFE, exkl. Short Positionierung)

Wie auch schon bei Strategie 1 sind erneut die a-posteriori Wahrscheinlichkeiten für einen steigenden oder fallenden Aktienindex, die mit Hilfe des Naive Bayes Klassifikators abgeleitet werden, die Basis der Entscheidungsfindung. Im Gegensatz zur eben ausgeführten simplen Strategie wird nicht dann in den Index investiert, wenn die a-posteriori Wahrscheinlichkeit für einen steigenden Index bei > 50 % liegt, sondern erst dann, wenn die a-posteriori Wahrscheinlichkeit den „Durchschnitt für Entscheidung“ („DFE“) übersteigt. In der anschließend dargestellten Tabelle 24 wird die Ermittlung des DFE dargestellt.

(1) Zeitpunkt der Investition t	(2) $P[DAX_{t+1} > DAX_t]$	(3) $P[DAX_{t+1} < DAX_t]$	(4) $P[DAX_{t+1} > DAX_t] > 50 \%$	(5) $P[DAX_{t+1} < DAX_t] > 50 \%$	(6) $\emptyset P[DAX_t > DAX_{t-1}]$	(7) $\emptyset P[DAX_t < DAX_{t-1}]$
15. Jänner 2017	61,91%	38,09%	61,91%	n.a.	69,39%	62,71%
15. Februar 2017	30,16%	69,84%	n.a.	69,84%	69,31%	62,71%
15. März 2017	81,56%	18,44%	81,56%	n.a.	69,31%	62,93%
15. April 2017	81,18%	18,82%	81,18%	n.a.	69,44%	62,93%
15. Mai 2017	50,97%	49,03%	50,97%	n.a.	69,57%	62,93%
15. Juni 2017	66,50%	33,50%	66,50%	n.a.	69,37%	62,93%
15. Juli 2017	73,61%	26,39%	73,61%	n.a.	69,34%	62,93%
15. August 2017	45,47%	54,53%	n.a.	54,53%	69,38%	62,93%
15. September 2017	62,13%	37,87%	62,13%	n.a.	69,38%	62,68%
15. Oktober 2017	83,43%	16,57%	83,43%	n.a.	69,31%	62,68%
15. November 2017	31,39%	68,61%	n.a.	68,61%	69,45%	62,68%
15. Dezember 2017	46,86%	53,14%	n.a.	53,14%	69,45%	62,85%

Tabelle 24: Ermittlung der DFE für 2017

In Spalte 1 ist der Zeitpunkt dargestellt, für den es eine Investitionsentscheidung zu treffen gilt. Die Spalten 2 und 3 zeigen die a-posteriori Wahrscheinlichkeit für einen steigenden respektive fallenden DAX und sind deckungsgleich mit den a-posteriori Wahrscheinlichkeiten bei Strategie 1, da sie auf die gleiche Art und Weise abgeleitet werden. In den Spalten 4 und 5 werden nur jene Wahrscheinlichkeiten für einen steigenden beziehungsweise fallenden Index dargestellt, die jeweils den Schwellwert in Höhe von 50 % überschreiten. Die Spalten 6 und 7 bilden den DFE sämtlicher Wahrscheinlichkeiten für einen steigenden beziehungsweise fallenden Index bis eine Periode vor dem Investitionszeitpunkt ab. Nur wenn die a-posteriori Wahrscheinlichkeiten den DFE übersteigen, ist potentiell eine Handlung erforderlich.

Für den Investitionszeitpunkt 15.02.2017 liegt die a-posteriori Wahrscheinlichkeit für einen fallenden Index bei rund 69,84 % (siehe Spalte 3). Da die a-posteriori Wahrscheinlichkeit in Höhe von rund 69,84 % den fraglichen DFE in Höhe von 62,71 % übersteigt, lautet die Handelsentscheidung ins Geld zu gehen, sofern sich die Strategie nicht bereits im Geld befindet. Dies ist in Tabelle 25 dargestellt.

(1) Zeitpunkt der Investition t	(2) $P[DAX_{t+1} > DAX_t]$	(3) $P[DAX_{t+1} < DAX_t]$	(4) $\emptyset P[DAX_t > DAX_{t-1}]$	(5) $\emptyset P[DAX_t < DAX_{t-1}]$	(6) Positionierung vor Handelsentscheidung	(7) Handelsentscheidung	(8) Positionierung nach Handelsentscheidung
15. Jänner 2017	61,91%	38,09%	69,39%	62,71%	im Geld	Keine Handlung	im Geld
15. Februar 2017	30,16%	69,84%	69,31%	62,71%	im Geld	Keine Handlung	im Geld
15. März 2017	81,56%	18,44%	69,31%	62,93%	im Geld	Long Instrument kaufen	Long im Markt
15. April 2017	81,18%	18,82%	69,44%	62,93%	Long im Markt	Keine Handlung	Long im Markt
15. Mai 2017	50,97%	49,03%	69,57%	62,93%	Long im Markt	Keine Handlung	Long im Markt
15. Juni 2017	66,50%	33,50%	69,37%	62,93%	Long im Markt	Keine Handlung	Long im Markt
15. Juli 2017	73,61%	26,39%	69,34%	62,93%	Long im Markt	Keine Handlung	Long im Markt
15. August 2017	45,47%	54,53%	69,38%	62,93%	Long im Markt	Keine Handlung	Long im Markt
15. September 2017	62,13%	37,87%	69,38%	62,68%	Long im Markt	Keine Handlung	Long im Markt
15. Oktober 2017	83,43%	16,57%	69,31%	62,68%	Long im Markt	Keine Handlung	Long im Markt
15. November 2017	31,39%	68,61%	69,45%	62,68%	Long im Markt	Long Instrument verkaufen	im Geld
15. Dezember 2017	46,86%	53,14%	69,45%	62,85%	im Geld	Keine Handlung	im Geld

Tabelle 25: Ermittlung der Handelsentscheidung nach Strategie 2 für 2017

Da Strategie 2 keine Positionen, in welchen man Short im Markt ist, vorsieht, beträgt die Rendite in jenen Monaten, in denen die Strategie im Geld ist, 0 %.

Die mit der Strategie 2 erzielten monatlichen Renditen werden in der folgenden Tabelle 26 beispielhaft für das Jahr 2017 dargestellt und mit der Benchmark verglichen.

(1) Zeitpunkt der Investition t	(2) Zeitpunkt der Rendite	(3) Positionierung im Aktienindex ab 15. des aktuellen Monats	(4) Rendite basierend auf Strategie 2	(5) DAX Rendite (ohne Strategie)	(6) Over-/Under-performance
15. Jänner 2017	15. Februar 2017	im Geld	0,00%	2,07%	-2,07%
15. Februar 2017	15. März 2017	im Geld	0,00%	1,83%	-1,83%
15. März 2017	18. April 2017	Long im Markt	-0,08%	-0,08%	0,00%
15. April 2017	15. Mai 2017	Long im Markt	6,72%	6,72%	0,00%
15. Mai 2017	15. Juni 2017	Long im Markt	-0,90%	-0,90%	0,00%
15. Juni 2017	17. Juli 2017	Long im Markt	-0,82%	-0,82%	0,00%
15. Juli 2017	15. August 2017	Long im Markt	-3,26%	-3,26%	0,00%
15. August 2017	15. September 2017	Long im Markt	2,81%	2,81%	0,00%
15. September 2017	16. Oktober 2017	Long im Markt	3,87%	3,87%	0,00%
15. Oktober 2017	15. November 2017	Long im Markt	-0,21%	-0,21%	0,00%
15. November 2017	15. Dezember 2017	im Geld	0,00%	0,98%	-0,98%
15. Dezember 2017	15. Jänner 2018	im Geld	0,00%	0,74%	-0,74%
Jahresrendite			8,05%	14,24%	

Tabelle 26: Darstellung der erzielten Renditen der Benchmark und nach Strategie 2 für 2017

Insgesamt wird bei Beibehaltung von Strategie 2 eine Rendite für 2017 in Höhe von 8,05 % erzielt. Bei einer Benchmark-Rendite von 14,24 % resultiert die Anwendung der Strategie 2 in einer erheblichen Underperformance, da die Strategie 2 für jene Monate, in denen der DAX positive Renditen erwirtschaftete, empfiehlt ins Geld zu gehen (etwa Jänner und Februar 2017). Für Monate wiederum, in denen der DAX gefallen ist, wurde von der Strategie 2 die Empfehlung Long im Markt zu sein, abgegeben.

7.3.3 Strategie 3 (DFE, inkl. Short Positionierung)

Folgende Tabelle 27 stellt die wesentlichen Entscheidungsparameter für Strategie 3 dar.

(1) Zeitpunkt der Investition t	(2) $P[DAX_{t+1} > DAX_t]$	(3) $P[DAX_{t+1} < DAX_t]$	(4) $\emptyset P[DAX_t > DAX_{t-1}]$	(5) $\emptyset P[DAX_t < DAX_{t-1}]$	(6) Positionierung vor Handelsentscheidung	(7) Handelsentscheidung	(8) Positionierung nach Handelsentscheidung
15. Jänner 2017	61,91%	38,09%	69,39%	62,71%	Short im Markt	Keine Handlung	Short im Markt
15. Februar 2017	30,16%	69,84%	69,31%	62,71%	Short im Markt	Keine Handlung	Short im Markt
15. März 2017	81,56%	18,44%	69,31%	62,93%	Short im Markt	Short Instrument verkaufen, Long Instrument kaufen	Long im Markt
15. April 2017	81,18%	18,82%	69,44%	62,93%	Long im Markt	Keine Handlung	Long im Markt
15. Mai 2017	50,97%	49,03%	69,57%	62,93%	Long im Markt	Keine Handlung	Long im Markt
15. Juni 2017	66,50%	33,50%	69,37%	62,93%	Long im Markt	Keine Handlung	Long im Markt
15. Juli 2017	73,61%	26,39%	69,34%	62,93%	Long im Markt	Keine Handlung	Long im Markt
15. August 2017	45,47%	54,53%	69,38%	62,93%	Long im Markt	Keine Handlung	Long im Markt
15. September 2017	62,13%	37,87%	69,38%	62,68%	Long im Markt	Keine Handlung	Long im Markt
15. Oktober 2017	83,43%	16,57%	69,31%	62,68%	Long im Markt	Keine Handlung	Long im Markt
15. November 2017	31,39%	68,61%	69,45%	62,68%	Long im Markt	Long Instrument verkaufen, Short Instrument kaufen	Short im Markt
15. Dezember 2017	46,86%	53,14%	69,45%	62,85%	Short im Markt	Keine Handlung	Short im Markt

Tabelle 27: Ermittlung der Handelsentscheidung nach Strategie 3 für 2017

In Spalte 1 ist der jeweilige Zeitpunkt der Investitionsentscheidung angegeben. Die Spalten 2 und 3 stellen die jeweiligen dazugehörigen a-posteriori Wahrscheinlichkeiten für einen erwarteten Anstieg beziehungsweise Rückgang des Aktienindex dar. Die Spalten 4 und 5 stellen die DFE sämtlicher Wahrscheinlichkeiten für einen steigenden beziehungsweise fallenden Index bis eine Periode vor dem Investitionszeitpunkt konsistent zu Tabelle 24 dar. Es werden wie in Unterabschnitt 7.3.2 bereits erwähnt nur Wahrscheinlichkeiten von größer als 50 % in die Durchschnittsbildung einbezogen. Da zum 15.01.2017 keiner der beiden Durchschnitte übertroffen wird, ist keine Handlung erforderlich. Es wird somit die Positionierung der Vorperiode (Short Position im Markt) beibehalten. Da hingegen zum 15.03.2017 die a-posteriori Wahrscheinlichkeit für einen steigenden Aktienindex in Spalte 2 größer als der DFE in Spalte 4 ist, wird das gehaltene Short Instrument veräußert und ein Long Instrument erworben.

Die mit der Strategie 3 erzielten monatlichen Renditen werden in der folgenden Tabelle 28 beispielhaft für das Jahr 2017 dargestellt und mit der Benchmark verglichen.

Zum 15.01.2017 empfiehlt die Strategie 3 eine Short Position im Markt. Da sich die Strategie bereits Short im Markt befindet, wird keine Handlung durchgeführt. Die bis zum 15.02.2017 generierte Rendite der Benchmark beträgt 2,07 %. Die mit Strategie 3 im selben Zeitraum generierte Rendite ergibt -2,07 % woraus sich eine Underperformance in Höhe von 4,14 % gegenüber der Benchmark ergibt. Über das gesamte Jahr betrachtet, ergibt sich unter Beibehaltung von Strategie 3 eine Rendite von 2,09 %.

(1) Zeitpunkt der Investition t	(2) Zeitpunkt der Rendite	(3) Positionierung im Aktienindex ab 15. des aktuellen Monats	(4) Rendite basierend auf Strategie 3	(5) DAX Rendite (ohne Strategie)	(6) Over-/Under-performance
15. Jänner 2017	15. Februar 2017	Short im Markt	-2,07%	2,07%	-4,14%
15. Februar 2017	15. März 2017	Short im Markt	-1,83%	1,83%	-3,66%
15. März 2017	18. April 2017	Long im Markt	-0,08%	-0,08%	0,00%
15. April 2017	15. Mai 2017	Long im Markt	6,72%	6,72%	0,00%
15. Mai 2017	15. Juni 2017	Long im Markt	-0,90%	-0,90%	0,00%
15. Juni 2017	17. Juli 2017	Long im Markt	-0,82%	-0,82%	0,00%
15. Juli 2017	15. August 2017	Long im Markt	-3,26%	-3,26%	0,00%
15. August 2017	15. September 2017	Long im Markt	2,81%	2,81%	0,00%
15. September 2017	16. Oktober 2017	Long im Markt	3,87%	3,87%	0,00%
15. Oktober 2017	15. November 2017	Long im Markt	-0,21%	-0,21%	0,00%
15. November 2017	15. Dezember 2017	Short im Markt	-0,98%	0,98%	-1,96%
15. Dezember 2017	15. Jänner 2018	Short im Markt	-0,74%	0,74%	-1,48%
Jahresrendite			2,09%	14,24%	

Tabelle 28: Darstellung der erzielten Renditen der Benchmark und nach Strategie 3 für 2017

7.3.4 Strategie 4 (HOE, keine Short Positionierung)

Wie bereits in Abschnitt 6.4 erwähnt, wird bei dieser Strategie mithilfe einer Datentabelle ermittelt, bei welchen kritischen Werten der a-posteriori Wahrscheinlichkeiten für einen steigenden oder fallenden Aktienindex die höchste Rendite im vorangegangenen Kalenderjahr erzielt worden wäre. Diese Analyse ist für das Jahr 2016 in nachfolgender Tabelle 29 dargestellt. Darin ist ersichtlich, dass bei einem kritischen Wert für „Aktienindex steigt“ von ≥ 50 % sowie für „Aktienindex fällt“ von ≥ 95 % eine Rendite von rund 21 % erzielt worden wäre.

	krit. Wert "Aktienindex fällt"										
	10%	50%	55%	60%	65%	70%	75%	80%	85%	90%	95%
krit. Wert "Aktienindex steigt"	50%	5%	5%	5%	15%	12%	10%	10%	10%	10%	21%
	55%	8%	8%	8%	18%	15%	10%	10%	10%	10%	21%
	60%	8%	8%	8%	18%	15%	10%	10%	10%	10%	21%
	65%	-4%	-4%	-4%	5%	5%	10%	10%	10%	10%	21%
	70%	-4%	-4%	-4%	5%	5%	10%	10%	10%	10%	21%
	75%	-4%	-4%	-4%	5%	5%	10%	10%	10%	10%	21%
	80%	-4%	-4%	-4%	5%	5%	10%	10%	10%	10%	21%
	85%	-4%	-4%	-4%	5%	5%	10%	10%	10%	10%	21%
	90%	-4%	-4%	-4%	5%	5%	10%	10%	10%	10%	21%
	95%	-4%	-4%	-4%	5%	5%	10%	10%	10%	10%	21%

Tabelle 29: Datentabelle, zur Ermittlung der kritischen Werte für 2017

Basierend auf den Ergebnissen in Tabelle 29 werden in einem nächsten Schritt in Tabelle 30 die jeweiligen a-posteriori Wahrscheinlichkeiten mit den kritischen Werten verglichen. Für den 15.01.2017 ergibt sich auf dessen Basis ein Signal eine Long-Position im Markt einzugehen, da die a-posteriori Wahrscheinlichkeit für DAX steigt (Spalte 2) mit 61,91 % im kritischen Bereich von 50,00 % - 100,00 % (Spalte 4) liegt. Da bereits eine Long-Position im Markt gehalten wird, ist folglich keine Handlung erforderlich.

(1) Zeitpunkt der Investition t	(2) $P[DAX_{t+1} > DAX_t]$	(3) $P[DAX_{t+1} < DAX_t]$	(4) Kritischer Bereich „Aktienindex steigt“	(5) Kritischer Bereich „Aktienindex fällt“	(6) Positionierung vor Handelsentscheidung	(7) Handelsentscheidung	(8) Positionierung nach Handelsentscheidung
15. Jänner 2017	61,91%	38,09%	50,00%-100,00%	95,00%-100,00%	Long im Markt	Keine Handlung	Long im Markt
15. Februar 2017	30,16%	69,84%	50,00%-100,00%	95,00%-100,00%	Long im Markt	Keine Handlung	Long im Markt
15. März 2017	81,56%	18,44%	50,00%-100,00%	95,00%-100,00%	Long im Markt	Keine Handlung	Long im Markt
15. April 2017	81,18%	18,82%	50,00%-100,00%	95,00%-100,00%	Long im Markt	Keine Handlung	Long im Markt
15. Mai 2017	50,97%	49,03%	50,00%-100,00%	95,00%-100,00%	Long im Markt	Keine Handlung	Long im Markt
15. Juni 2017	66,50%	33,50%	50,00%-100,00%	95,00%-100,00%	Long im Markt	Keine Handlung	Long im Markt
15. Juli 2017	73,61%	26,39%	50,00%-100,00%	95,00%-100,00%	Long im Markt	Keine Handlung	Long im Markt
15. August 2017	45,47%	54,53%	50,00%-100,00%	95,00%-100,00%	Long im Markt	Keine Handlung	Long im Markt
15. September 2017	62,13%	37,87%	50,00%-100,00%	95,00%-100,00%	Long im Markt	Keine Handlung	Long im Markt
15. Oktober 2017	83,43%	16,57%	50,00%-100,00%	95,00%-100,00%	Long im Markt	Keine Handlung	Long im Markt
15. November 2017	31,39%	68,61%	50,00%-100,00%	95,00%-100,00%	Long im Markt	Keine Handlung	Long im Markt
15. Dezember 2017	46,86%	53,14%	50,00%-100,00%	95,00%-100,00%	Long im Markt	Keine Handlung	Long im Markt

Tabelle 30: Handlungsempfehlung nach Strategie 4 für 2017

Basierend auf den in Tabelle 30 dargestellten Positionierungen ist in Tabelle 31 die mit Strategie 4 erzielte Rendite für das Jahr 2017 dargestellt. Diese beträgt 14,24 % und ist deckungsgleich mit jener der Benchmark DAX.

(1) Zeitpunkt der Investition t	(2) Zeitpunkt der Rendite	(3) Positionierung im Aktienindex ab 15. des aktuellen Monats	(4) Rendite basierend auf Strategie 4	(5) DAX Rendite (ohne Strategie)	(6) Over-/Under-performance
15. Jänner 2017	15. Februar 2017	Long im Markt	2,07%	2,07%	0,00%
15. Februar 2017	15. März 2017	Long im Markt	1,83%	1,83%	0,00%
15. März 2017	18. April 2017	Long im Markt	-0,08%	-0,08%	0,00%
15. April 2017	15. Mai 2017	Long im Markt	6,72%	6,72%	0,00%
15. Mai 2017	15. Juni 2017	Long im Markt	-0,90%	-0,90%	0,00%
15. Juni 2017	17. Juli 2017	Long im Markt	-0,82%	-0,82%	0,00%
15. Juli 2017	15. August 2017	Long im Markt	-3,26%	-3,26%	0,00%
15. August 2017	15. September 2017	Long im Markt	2,81%	2,81%	0,00%
15. September 2017	16. Oktober 2017	Long im Markt	3,87%	3,87%	0,00%
15. Oktober 2017	15. November 2017	Long im Markt	-0,21%	-0,21%	0,00%
15. November 2017	15. Dezember 2017	Long im Markt	0,98%	0,98%	0,00%
15. Dezember 2017	15. Jänner 2018	Long im Markt	0,74%	0,74%	0,00%
Jahresrendite			14,24%	14,24%	

Tabelle 31: Darstellung der erzielten Renditen der Benchmark und nach Strategie 4 für 2017

7.4 Gesamtrendite und Vergleich mit der Benchmark

Die in den Abschnitten 7.1 bis 7.3 dargestellten Analysen wurden für die Jahre 2004 bis 2019 sowohl für den deutschen (DAX) als auch für den amerikanischen Markt (S&P 500) durchgeführt. Im Folgenden werden für beide Märkte die Gesamtrenditen der vier Handelsstrategien dargestellt und mit der Rendite der Benchmark verglichen. Es wird auf etwaige Gemeinsamkeiten der einzelnen Strategien eingegangen und ein Resümee gezogen.

7.4.1 Deutscher Markt – DAX

Die für den deutschen Markt resultierenden jährlichen prozentualen Renditen des Buy-and-Hold Portfolios (Benchmark) sowie der Portfolios basierend auf den vier Handelsstrategien sind in Abbildung 23 dargestellt:

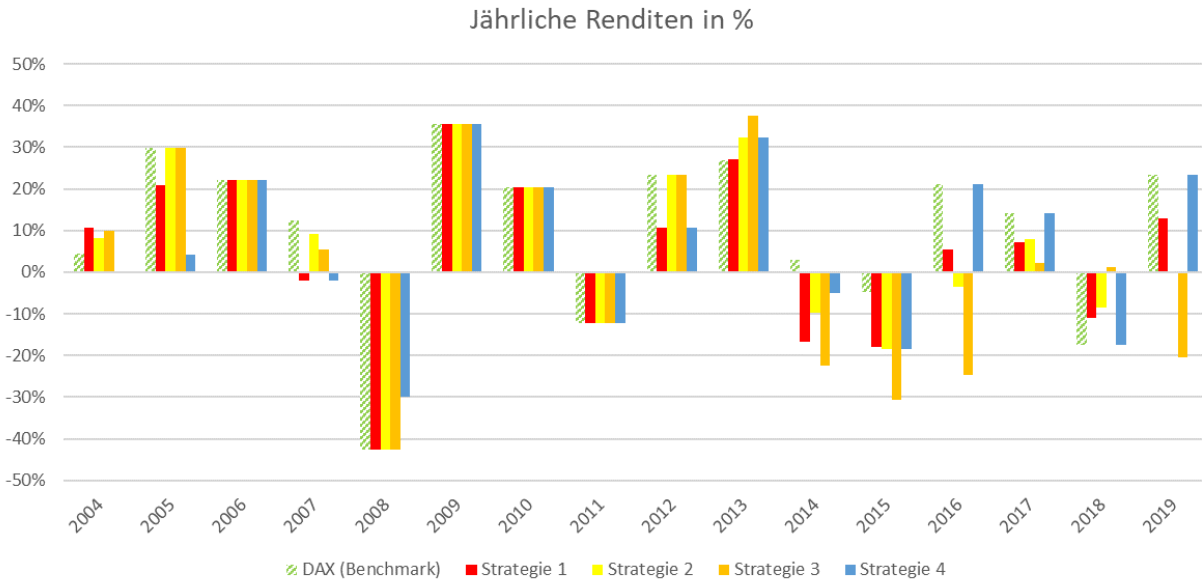


Abbildung 23: Jährliche Renditen im Vergleich - DAX

Die Gesamtpformance von 2004 bis 2019 für die Strategien 1 bis 3 sowie von 2005 bis 2019 für Strategie 4 jeweils im Vergleich zur Benchmark sowie die durchschnittliche jährliche Rendite je Portfolio und Benchmark sind in nachfolgender Tabelle 32 dargestellt:

Strategie	Perioden	Ø Jahresrendite je Strategie	Ø Jahresrendite Benchmark	Performance je Strategie	Performance Benchmark DAX	Underperformance in %-punkten (PP = Prozentpunkte)
Strategie 1	2004-2019	2,71%	14,38%	43,43%	230,13%	-186,71 PP
Strategie 2	2004-2019	4,80%	14,38%	76,87%	230,13%	-153,26 PP
Strategie 3	2004-2019	-0,92%	14,38%	-14,70%	230,13%	-244,84 PP
Strategie 4	2005-2019	6,69%	14,43%	100,33%	216,39%	-116,06 PP

Tabelle 32: Gegenüberstellung der Performance je Strategie und DAX

Es ist deutlich zu erkennen, dass es mit keiner der untersuchten Strategien gelang, im Beobachtungszeitraum eine einfache Buy-and-Hold Strategie der Benchmark DAX zu übertreffen.

Eine detaillierte Übersicht über die jährlich erzielte Out- beziehungsweise Underperformance der unterschiedlichen Handelsstrategien, ist nachfolgend in Tabelle 33 dargestellt:

Jährliche Out- / Underperformance	Strategie 1	Strategie 2	Strategie 3	Strategie 4
2004	6,37 PP	3,95 PP	5,65 PP	n.a.
2005	-9,00 PP	0,00 PP	0,00 PP	-25,72 PP
2006	0,00 PP	0,00 PP	0,00 PP	0,00 PP
2007	-14,48 PP	-3,29 PP	-7,06 PP	-14,48 PP
2008	0,00 PP	0,00 PP	0,00 PP	12,83 PP
2009	0,00 PP	0,00 PP	0,00 PP	0,00 PP
2010	0,00 PP	0,00 PP	0,00 PP	0,00 PP
2011	0,00 PP	0,00 PP	0,00 PP	0,00 PP
2012	-12,66 PP	0,00 PP	0,00 PP	-12,66 PP
2013	0,18 PP	5,43 PP	10,86 PP	5,43 PP
2014	-19,86 PP	-12,84 PP	-25,57 PP	-8,06 PP
2015	-13,14 PP	-13,47 PP	-25,88 PP	-13,47 PP
2016	-15,59 PP	-24,60 PP	-45,67 PP	0,00 PP
2017	-6,97 PP	-6,20 PP	-12,15 PP	0,00 PP
2018	6,50 PP	9,13 PP	18,77 PP	0,00 PP
2019	-10,47 PP	-23,32 PP	-43,88 PP	0,00 PP

Tabelle 33: Jährliche Out- / Underperformance je Strategie - DAX

Nachfolgende Abbildung 24 veranschaulicht die Entwicklung eines Portfolios mit einem Startwert von 100 Geldeinheiten am 15.01.2004 für die angewandten Strategien 1 bis 3. Die Strategie 4 startet aufgrund ihrer Spezifika ein Jahr später (blaue Linien). Zusätzlich ist die Entwicklung eines Portfolios basierend auf der Buy-and-Hold Strategie in der Benchmark dargestellt (orangene Linien). Die vertikalen grauen Balken stellen die jeweilige Out- oder Underperformance der Strategien gegenüber der Benchmark in Prozent dar.

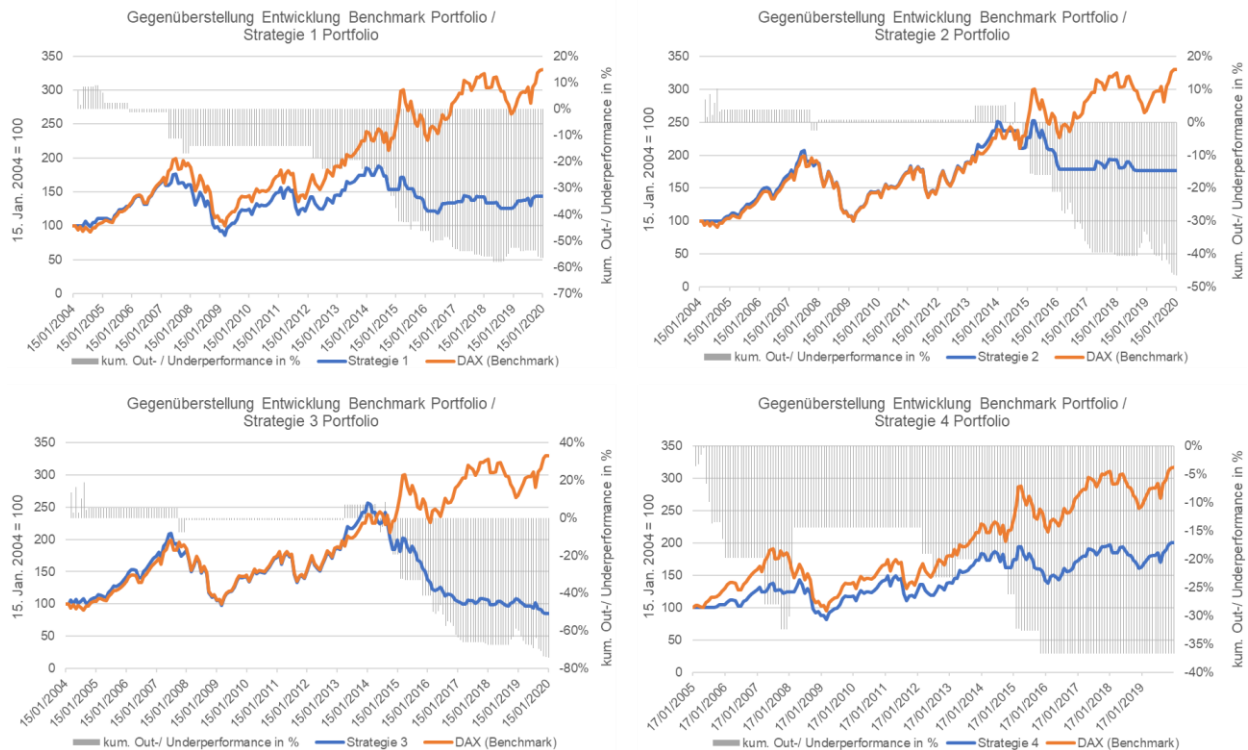


Abbildung 24: Entwicklung Portfoliowert je Strategie - DAX

Während sich die Portfolios basierend auf den Strategien 2 und 3 in etwa bis Mitte 2014 relativ ident zum Benchmark-Portfolio entwickelten, erzielten die Portfolios basierend auf Strategien 1 (nach einer kurzzeitigen Outperformance) und 4 bereits sehr früh eine Underperformance gegenüber dem Benchmark-Portfolio. Abbildung 23 stellt diesen Sachverhalt anschaulich dar. Es ist zu erkennen, dass bis zum Jahr 2014, die jährlichen Renditen der Strategien 2 und 3 in etwa jenen der Benchmark entsprechen. Identische Entwicklungen der jährlichen Renditen zwischen den einzelnen Strategien sowie der Benchmark lassen sich beispielsweise im Jahr 2006 im Wesentlichen darauf zurückzuführen, dass alle vier Strategien für das gesamte Jahr von steigenden Kursen ausgingen, dementsprechend waren alle vier Portfolios gleichzeitig mit einer Long-Position investiert. In den Jahren 2009, 2010 und 2011 lässt sich die idente Entwicklung dadurch erklären, dass die Short Listen für diese Investitionsjahre keine Indikatoren inkludierten, da wie in Abschnitt 4.6 gefordert, kein Indikator simultan sowohl nach der Forward Selection Methode als auch nach der Backward Elimination Methode das statistische Signifikanzkriterium erfüllte. Um trotzdem in diesen Jahren eine Investitionsentscheidung treffen zu können, wurde für diese Jahre jene für

Dezember 2008 getroffene Entscheidung, eine Kaufentscheidung, beibehalten, bis im Jänner 2012, basierend auf der abgeleiteten Short Liste der Untersuchungsjahre 2008 bis 2011 eine neuerliche Investitionsentscheidung getroffen werden konnte.

Schlussendlich zeigt sich Ende 2019 die beste Entwicklung der vier Portfolios im Portfolio der Strategie 4 mit einer Gesamtrendite von 100,33 %, gefolgt von Strategie 2 mit einer Gesamtrendite in Höhe von 76,87 %. Ungeachtet der immer noch deutlichen Underperformance gegenüber der Benchmark lässt die relativ „bessere“ Entwicklung des Portfolios der Strategie 4 darauf deuten, dass eine Entscheidung auf Basis des kritischen Werts der a-posteriori Wahrscheinlichkeit, welcher im vorangegangenen Kalenderjahr zur besten Performance geführt hatte, als durchaus sinnvoll erachtet werden kann.

7.4.2 Amerikanischer Markt – S&P 500

Die jährlichen Renditen, die mit den im Rahmen der vorliegenden Arbeit elaborierten Strategien erzielt wurden, werden über den gesamten Investitionszeitraum jenen Renditen gegenübergestellt, die der S&P 500 im selben Zeitraum erzielt. Details zu den Jahresrenditen können der nachfolgenden Abbildung 25 entnommen werden.

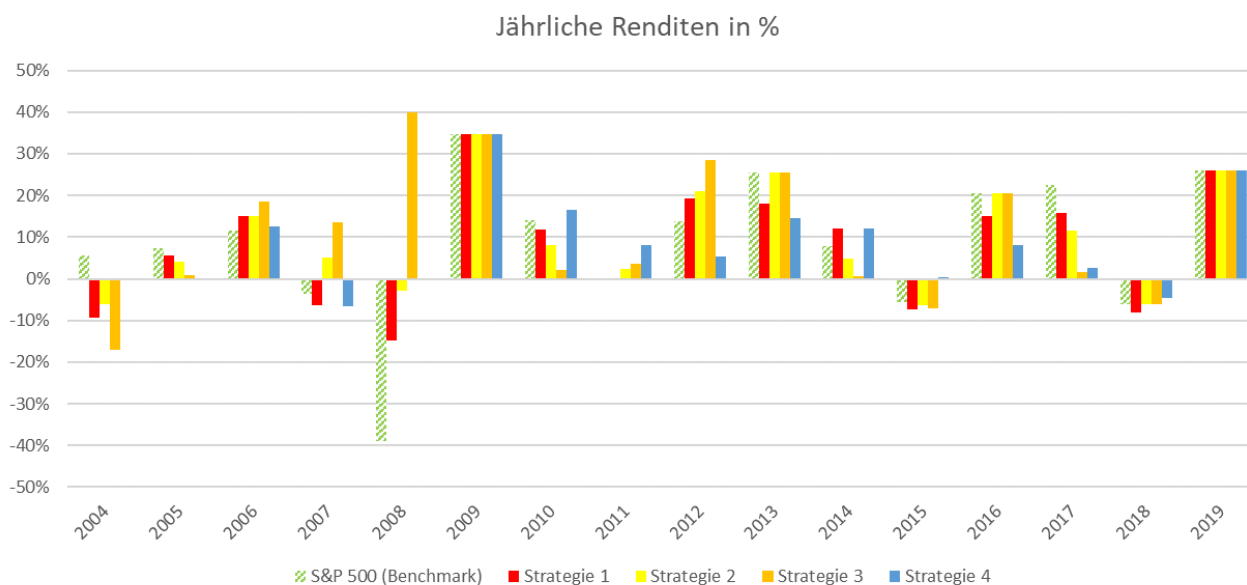


Abbildung 25: Jährliche Renditen im Vergleich – S&P 500

Die nachfolgende Tabelle 34 zeigt die durchschnittliche Jahresrendite je Strategie und der Benchmark S&P 500 sowie die Gesamtrenditen, die mit der jeweiligen

Investitionsstrategie beziehungsweise dem S&P 500 erzielt wurden. Darüberhinausgehend wird die Outperformance, die sich aus der Gegenüberstellung mit den Gesamtrenditen ergibt, in Prozentpunkten („PP“) gezeigt.

Alle Investitionsstrategien haben gemeinsam, dass sie zu einer Outperformance im Vergleich zur gewählten Benchmark führen, wobei die Outperformance, die mit der simplen Strategie erzielt wird, sehr gering ausfällt. Strategie 3 hingegen ist jene Strategie, die den höchsten Vermögenszuwachs erwirtschaftet und retrospektiv die präferierte Entscheidungsgrundlage sein sollte.

Strategie	Perioden	Ø Jahresrendite je Strategie	Ø Jahresrendite Benchmark	Performance je Strategie	Performance Benchmark S&P500	Outperformance in %-punkten
Strategie 1	2004-2019	12,42%	11,91%	198,69%	190,56%	8,13 PP
Strategie 2	2004-2019	19,22%	11,91%	307,52%	190,56%	116,96 PP
Strategie 3	2004-2019	24,39%	11,91%	390,31%	190,56%	199,75 PP
Strategie 4	2005-2019	14,89%	11,67%	223,37%	175,03%	48,35 PP

Tabelle 34: Gegenüberstellung der Performance je Strategie und S&P 500

Eine detaillierte Übersicht über die jährlich erzielte Out- beziehungsweise Underperformance der unterschiedlichen Handelsstrategien, ist nachfolgend in Tabelle 35 dargestellt:

Jährliche Out-/ Underperformance	Strategie 1	Strategie 2	Strategie 3	Strategie 4
2004	-15,08 PP	-11,63 PP	-22,58 PP	NA
2005	-1,59 PP	-3,13 PP	-6,49 PP	-7,27 PP
2006	3,41 PP	3,41 PP	6,81 PP	1,03 PP
2007	-2,67 PP	8,65 PP	17,10 PP	-2,92 PP
2008	24,17 PP	36,04 PP	78,94 PP	38,90 PP
2009	0,00 PP	0,00 PP	0,00 PP	0,00 PP
2010	-2,16 PP	-5,91 PP	-11,96 PP	2,58 PP
2011	-0,28 PP	2,35 PP	3,75 PP	8,15 PP
2012	5,59 PP	7,30 PP	14,60 PP	-8,47 PP
2013	-7,38 PP	0,00 PP	0,00 PP	-11,07 PP
2014	4,33 PP	-2,97 PP	-7,31 PP	4,33 PP
2015	-1,74 PP	-0,72 PP	-1,43 PP	6,13 PP
2016	-5,59 PP	0,00 PP	0,00 PP	-12,45 PP
2017	-6,54 PP	-10,74 PP	-20,89 PP	-19,85 PP
2018	-2,12 PP	0,00 PP	0,00 PP	1,27 PP
2019	0,00 PP	0,00 PP	0,00 PP	0,00 PP

Tabelle 35: Jährliche Out- / Underperformance je Strategie - S&P 500

In weiterer Folge werden die Entwicklungen des Vergleichsindex und der Renditen, die mit den unterschiedlichen Handelsstrategien erzielt werden, graphisch veranschaulicht:

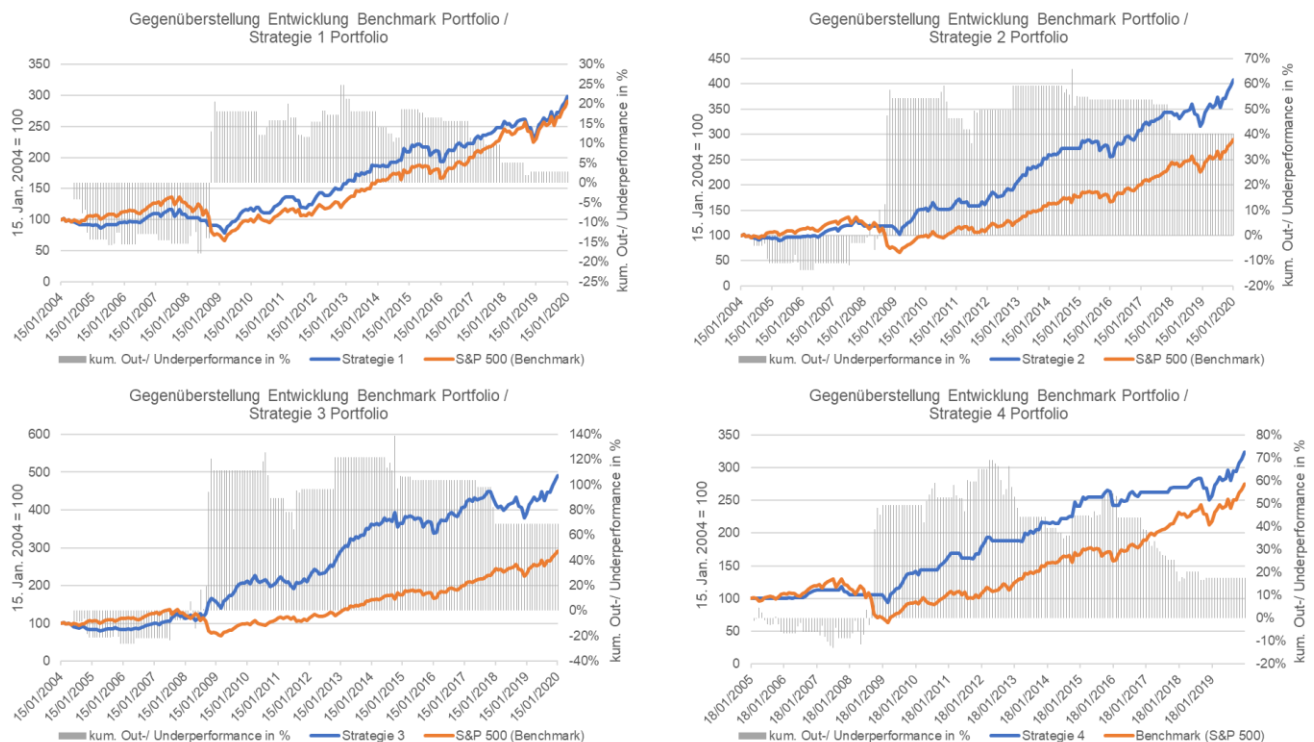


Abbildung 26: Entwicklung Portfoliowert je Strategie – S&P 500

Die blauen und orangenen Linien sowie die grauen Balken sind analog zur Erläuterung in Unterabschnitt 7.4.1 zu verstehen.

Der Verlauf der Rendite der simplen Strategie ähnelt über den gesamten Betrachtungszeitraum stark der Entwicklung der Benchmark. Liegen die kumulierten Gewinne, die mit Strategie 1 erzielt werden, zu Beginn noch konstant unter jenen des S&P 500, kann im Oktober 2008 erstmals die kumulierte Rendite des Vergleichsindex übertroffen und in den Folgejahren die Outperformance noch weiter ausgebaut werden. Bis zum Ende des Betrachtungszeitraums wird mit Strategie 1 eine Outperformance erzielt, wobei sich diese gegen Ende des Betrachtungszeitraums abbaut.

Wie dies auch bei der eben beschriebenen Strategie 1 der Fall ist, wird ebenso mit Strategie 2 in den ersten Jahren eine geringfügige Underperformance erzielt. Erneut markiert das Jahr 2008 einen Wendepunkt und ein Investor, der sich Strategie 2 bedient, kann sich erstmals besserstellen. Die Outperformance wird im Laufe des

Investitionszeitraums weiter ausgebaut, so dass der blaue Liniengraph deutlich über jenem Liniengraph liegt, der den S&P 500 repräsentiert.

Die Entwicklung der Rendite, die mit Strategie 3 erzielt wird, ähnelt dem Verlauf der Gewinne der zweiten Strategie: Erneut bleiben zu Beginn des Investitionszeitraums die mit der entwickelten Handelsstrategie erwirtschafteten kumulierten Gewinne hinter jenen der einfachen Buy-and-Hold Strategie zurück, und abermals wird im Jahr 2008 das erste Mal die Benchmark übertroffen. Im Unterschied zu den anderen zuvor erläuterten Strategien, kann die Outperformance wesentlich stärker ausgebaut werden. Besonders auffällig ist der sprunghafte Anstieg der Outperformance im Jahr 2013.

Das Portfolio basierend auf Strategie 4 entwickelt sich in den ersten Investitionsjahren ähnlich wie die bereits erläuterten Strategien. Im Wesentlichen zeigt ein Vergleich der beiden Liniengraphen, dass die Benchmark- und Strategie 4-Rendite bis zum Jahr 2008 beinahe einen identen Verlauf aufweisen, wobei die kumulierte Rendite nach Strategie 4 kontinuierlich unter jener der Benchmark liegt. Wie auch schon bei den Strategien zuvor, ändert sich erneut ab dem Jahr 2008 die Effektivität der Handelsstrategie. Es wird eine höhere Outperformance erzielt, die in weiterer Folge ausgebaut wird. Die Outperformance erreicht Mitte 2012 ihren Höhepunkt mit rund 69 %. Ab diesem Zeitpunkt entwickelt sie sich rückläufig und nimmt ab Anfang 2018 einen konstanten Verlauf, da sich die Rendite nach Strategie 4 ident zur Benchmark-Rendite entwickelt.

Da die entwickelten Investitionsstrategien allesamt auf den a-posteriori Wahrscheinlichkeiten basieren, weisen die unterschiedlichen Portfolios vereinzelt deckungsgleiche Jahresrenditen auf. Zwei Jahre, die in diesem Zusammenhang besonders hervorstechen, sind 2009 und 2019, da sich die Portfolios der Strategien 1 bis 4 ident mit der Benchmark entwickeln. Dieser Umstand ist dadurch begründet, dass die Handelsempfehlungen sämtlicher Strategien über die beiden Jahre hinweg gesamthaft auf Long im Markt stehen. Daraus resultieren Portfoliorenditen, die sich ident zur Benchmark entwickeln.

7.4.3 Gegenüberstellung der Effektivität der Handelsstrategien in den Vergleichsmärkten

Während alle Strategien am amerikanischen Aktienmarkt zu einer Outperformance führen, sind die dargestellten Strategien am deutschen Markt wenig effektiv – es gelingt in keinem Szenario die Rendite des DAX zu übertreffen. Auch hinsichtlich der Profitabilität der entwickelten Strategien zeigen die beiden Märkte keine Gemeinsamkeiten. Aus den vorangegangenen Unterabschnitten 7.4.1 und 7.4.2 ist ersichtlich, dass weder die erfolgreichste noch die am wenigsten erfolgreiche Strategie für die beiden Märkte identisch ist. Mögliche Erklärungsversuche für die abweichenden Ergebnisse in den beiden Märkten könnten auf die im nachfolgenden Kapitel dargestellten Limitationen zurückzuführen sein.

Basierend auf den Ergebnissen könnte für den deutschen Markt davon ausgegangen werden, dass die EMH zumindest in ihrer halbstrengen Form Gültigkeit besitzt. Diese Erkenntnisse stehen in Einklang mit Plíhal (2016). Plíhal (2016, S. 328) testet auf Granger-Kausalität zwischen dem DAX und ausgewählten volkswirtschaftlichen Indikatoren wie beispielsweise der industriellen Produktion oder der Handelsbilanz im Zeitraum zwischen Jänner 1999 und September 2015. Der Autor kommt zu dem Schluss, dass kein untersuchter makroökonomischer Indikator den DAX erklären kann (anhand Granger-Kausalität). Somit kann laut Plíhal (2016, S. 328) darauf geschlossen werden, dass der deutsche Aktienmarkt informationseffizient ist und nicht gegen die EMH verstößt.

8 Limitationen

Die Ergebnisse der vorgestellten Handelsstrategien sollten unter Kenntnis der, in diesem Abschnitt dargestellten, Limitationen bewertet werden. Diese erstrecken sich von der Auswahl der Indikatoren bis hin zur Auswertung der Handelsergebnisse je Strategie.

Die Indikatoren, die in dieser Arbeit als Inputvariablen für die Klassifikation dienen, beziehen sich nur auf volkswirtschaftliche Tatbestände. Es werden keinerlei technische Indikatoren (beispielsweise gleitende Durchschnitte) oder vergangene Kursbewegungen herangezogen. Des Weiteren wurde zwar eine systematische Recherche für die

Erstellung der Long List an volkswirtschaftlichen Indikatoren durchgeführt, diese kann allerdings nicht als allumfassend angesehen werden. So wurden beispielsweise die Indikatoren anhand deren „Wichtigkeit“ auf Basis der drei genannten Quellen in Abschnitt 5.2 abgeleitet, allerdings wurde dabei eine Stichtagsbetrachtung durchgeführt. Veränderungen der „Wichtigkeit“ im Zeitverlauf wurden somit in der vorliegenden Arbeit nicht berücksichtigt. Die Auswahl abweichender Quellen kann somit zu anderen Long Lists für die betrachteten Märkte und somit auch zu einer abweichenden Auswahl von Indikatoren als Inputvariablen für die Klassifikation führen.

Des Weiteren wurden im Zuge der Überführung der Long List in die Short List nur jene Indikatoren herangezogen, die sowohl in der Forward Regression als auch in der Backward Elimination Methodik simultan eine signifikante Erklärungsfähigkeit aufweisen konnten (definiert anhand $\alpha=0,1$). Für die Durchführung dieser zwei Regressionen wurde die monatliche Rendite des DAX (S&P 500), der abhängigen Variablen, zu einem bestimmten Zeitpunkt betrachtet. Da die Indikatoren, welche in die Regression als unabhängige Variablen einfließen, in der Mehrzahl der Fälle an unterschiedlichen Tagen des Monats veröffentlicht werden, war es nicht möglich in der Regression idente Zeitabstände (in Tagen) zwischen der Rendite des DAX (S&P 500) und des jeweiligen Datenpunkts eines jeden Indikators sicherzustellen.

In Abschnitt 3.2 wurde eine detaillierte Analyse vorgestellt, die das Rational zur Auswahl des Naive Bayes Indikators zur Generierung der Handelssignale liefert. Wenngleich dieser für die vorliegende Arbeit als beste Alternative unter den vorgestellten Klassifikatoren gewählt wurde, gibt es allerdings noch zahlreiche andere Klassifikatoren, die in der vorliegenden Arbeit nicht näher betrachtet wurden. Dies hat teils unterschiedliche Gründe. Ein wesentlicher Eckpfeiler dieser Arbeit ist es, wie bereits erwähnt, eine einfach umsetzbare Handelsstrategie zu etablieren, wobei die ausgewählten Klassifikatoren unseres Erachtens diesem Anspruch am besten gerecht werden.

Eine weitere Limitation richtet sich auf die Ausgestaltung des Naive Bayes Klassifikators. Betrachtet man diesen Klassifikator unter Verwendung kontinuierlicher Daten, könnte

man beispielsweise auch andere Verteilungen als die Normalverteilung zur Ermittlung der Likelihoods zugrunde legen.

Des Weiteren muss, unabhängig davon, ob die Strategien in der Lage sind die Benchmark zu schlagen oder nicht, beachtet werden, dass die Renditen aller Portfolios ohne Berücksichtigung etwaiger Transaktionskosten, Steuern und eines vorherrschenden Bid-Ask Spreads ermittelt wurden. Eine Berücksichtigung dieser Komponenten würde zu einer Verschlechterung der Rendite der Portfolios der vier Strategien je Markt führen, da die Handelsfrequenz hier deutlich höher ist als bei dem Benchmark Portfolio und diese Kosten im Zuge jedes Handelsablaufes anfallen. Bei den Portfolios basierend auf den vier Strategien wird immer dann gehandelt, wenn der Klassifikator ein Handelssignal generiert, beim Benchmark Portfolio hingegen wird nur zu Beginn der Untersuchungsperiode (Anfang 2004) gekauft und am Ende (Anfang 2020) verkauft.

Schlussendlich stellt die in dieser Arbeit dargestellte Analyse eine Vergangenheitsbetrachtung dar. Dementsprechend müssen die Erkenntnisse nicht zwangsweise auch für zukünftige Perioden gelten und es kann somit von den Ergebnissen der Betrachtungsperiode nicht auf einen Erfolg oder Misserfolg der Strategien in zukünftigen Perioden geschlossen werden.

9 Zusammenfassung und Ausblick

Ziel der vorliegenden Masterarbeit war es eine Handelsstrategie zu kreieren mit deren Hilfe eine Outperformance im Vergleich zu den gewählten Benchmarks, dem DAX sowie dem S&P 500, erreicht werden kann.

Da die Handelsstrategie auf intuitiv verständlichen, objektiven und zudem rechnerisch wenig aufwendigen Kriterien basieren sollte, wurde der Naive Bayes Klassifikator herangezogen. Wesentlicher Vorteil des Naive Bayes Klassifikators gegenüber anderen Klassifikatoren, die im Rahmen der Arbeit beleuchtet wurden, ist der Umstand, dass sich dieser gleichermaßen für die Anwendung bei einer geringen als auch einer großen Anzahl an Attributen eignet – der rechnerische Aufwand ist von der Anzahl der Attribute relativ unabhängig (Aggarwal, 2015, S. 67). Neben dem Naive Bayes Klassifikator dienten makroökonomische Indikatoren, ob ihrer in der Literatur vielfach diskutierten Interdependenzen mit dem Aktienmarkt als wesentlicher Eckpfeiler der entwickelten Handelsstrategie. Welche der volkswirtschaftlichen Indikatoren dabei maßgeblich für die Investitionsentscheidung waren, war von statistischen Auswertungen einer Grundgesamtheit an volkswirtschaftlichen Indikatoren abhängig: Nach einer Überprüfung der Indikatoren auf deren Stationarität, wurden diese auf das Vorliegen von Multikollinearität getestet. Um möglichst präzise Ergebnisse zu erhalten, fanden schlussendlich nur jene Indikatoren Eingang in den Entscheidungsprozess, die im schrittweisen Selektionsprozess sowohl nach dem Verfahren der „Forward Selection“ als auch nach dem Verfahren der „Backward Elimination“ selektiert wurden.

Die so selektierten volkswirtschaftlichen Daten bildeten die Ausgangslage für die Ableitung der a-posteriori Wahrscheinlichkeiten, welche wiederum maßgeblich für die Investitionsentscheidungen nach dem so genannten „Naive Bayes Investing“ waren. Abhängig davon, welchen Schwellwert die a-posteriori Wahrscheinlichkeit übersteigen musste (etwa $> 50\%$), um eine Verkaufs- und Kaufhandlung auszulösen, wurden vier unterschiedliche Strategien abgeleitet, die allerdings allesamt auf den Naive Bayes a-posteriori Wahrscheinlichkeiten fußen. Neben den verschiedenen Schwellwerten unterscheiden sich die Strategien zudem dahingehend, ob Investoren im Rahmen der

jeweiligen Strategie auch von fallenden Aktienkursen profitieren, also eine Short-Positionierung im Markt eingehen können.

Zusammenfassend sei an dieser Stelle festgehalten, dass es mit dem „Naive Bayes Investing“ gelungen ist, eine Handelsstrategie zu entwickeln, die Objektivität bei der Investitionsentscheidung sowie eine relativ einfache Anwendbarkeit mit einer simplen Beschaffbarkeit der zugrundeliegenden Datenbasis vereint. Hinsichtlich des Erfolgs der Strategie im Hinblick auf die Erzielung von Überrenditen im Vergleich zu den Referenzmärkten zeigen sich die Ergebnisse allerdings divergent. Während es beim DAX mit keiner Strategie gelungen ist sich besser als der Markt zu stellen, zeigt sich beim S&P 500 ein gänzlich anderes Bild: Investoren, die Anlageentscheidungen auf Basis der dargestellten Strategien treffen, erzielen mit allen vier Strategien eine bessere Rendite als mit einer einfachen Buy-and-Hold Strategie des S&P 500. Eine mögliche Ursache für die mangelnde Effektivität der Investmentstrategie beim DAX liefert die Publikation von Plíhal (2016), die nahe legt, dass der deutsche Aktienmarkt informationseffizient ist und die EMH somit zumindest in ihrer halbstrengen Form Gültigkeit besitzt.

Wenngleich die elaborierten Strategien am deutschen Aktienmarkt nicht die angestrebten Ergebnisse erzielten, so soll an dieser Stelle nochmal hervorgehoben werden, dass die Strategien für beide Referenzmärkte nur eine Vergangenheitsbetrachtung darstellen und keinerlei Rückschlüsse über die Effektivität des „Naive Bayes Investing“ für künftige Perioden zulassen. Weitere Einschränkungen in Zusammenhang mit den entwickelten Strategien erheben sich daraus, dass die Auswahl der Long List an volkswirtschaftlichen Indikatoren nicht als allumfassend angesehen werden kann. Zudem wurde für jeden betrachteten Aktienmarkt eine gesonderte Long List erstellt. Dieser Umstand macht das „Naive Bayes Investing“, wie es in dieser Arbeit umgesetzt wurde, wenig flexibel in dem Sinne, dass für jeden Aktienmarkt im Vorfeld zunächst erhebliche zeitliche Ressourcen erforderlich sind, um jene Indikatoren auszuwählen, die potentiell Relevanz für den jeweiligen Markt haben könnten. Ein weiterer Kritikpunkt stellt die Einschränkung der Inputvariable auf volkswirtschaftliche Indikatoren dar. Weder technische Indikatoren noch vergangene Kursbewegungen finden Eingang in die Analyse.

Der Versuch ein einheitliches Set an Inputvariablen zu finden, das eine Kombination aus makroökonomischen und technischen Indikatoren gleichermaßen berücksichtigt, und auf eine Vielzahl von Märkten anwendbar ist, könnten Ausgangspunkt für weitere Analysen bilden. Auch die divergierende Wirksamkeit des „Naive Bayes Investing“ in den beiden analysierten Märkten (DAX und S&P 500) könnten weiter untersucht werden und möglicherweise darüber Aufschluss geben, wovon der Erfolg dieser Strategie abhängt.

Naive Bayes Klassifikatoren stellen jedenfalls aufgrund ihrer dargelegten Vorteile ein wichtiges Instrument dar, das bereits in den unterschiedlichsten Fachbereichen Eingang gefunden hat und erfolgreich eingesetzt werden konnte. Eine Anwendung dieser Klassifikatoren auf weitere Forschungsbereiche scheint somit wahrscheinlich.

References

- Abu Alfeilat, H. A., Hassanat, A. B. A., Lasassmeh, O., Tarawneh, A. S., Alhasanat, M. B., Eyal Salman, H. S. and Prasath, V. B. S. (2019) 'Effects of Distance Measure Choice on K-Nearest Neighbor Classifier Performance: A Review', *Big data*, vol. 7, no. 4, pp. 221–248.
- Aggarwal, C. C. (ed) (2015) *Data classification: Algorithms and applications / edited by Charu C. Aggarwal*, Boca Raton, Chapman & Hall/CRC.
- Anderson, D. R. (2020) *Modern business statistics with Microsoft Excel*, Mason, South-Western.
- Arsad, Z. b., Aik, C. S. and Nordin, S. N. M. (2011) 'Predictability of Turn-of-the-Month Effect at Stock Markets in Malaysia, South Korea and Japan', *SSRN Electronic Journal* [Online]. DOI: 10.2139/ssrn.1913882.
- Asprem, M. (1989) 'Stock prices, asset portfolios and macroeconomic variables in ten European countries', *Journal of Banking & Finance*, vol. 13, 4-5, pp. 589–612.
- Baumohl, B. (2012) *The secrets of economic indicators: Hidden clues to future economic trends and investment opportunities / Bernard Baumohl*, 3rd edn, Upper Saddle River, N.J., FT Press.
- Black, K. (2019) *Business statistics: For contemporary decision making*, Hoboken, Wiley.
- Bloomberg L.P. (2020) *Bloomberg Terminal* [Computer program]. Available at <https://www.bloomberg.com/professional/solution/bloomberg-terminal/> (Accessed 10 October 2020).
- Boersenlexikon.faz.net. (ed) (2021) *Aktienindex: Definition im FAZ.NET Börsenlexikon*. [Online]. Available at <https://boersenlexikon.faz.net/definition/aktienindex/> (Accessed 29 May 2021).
- börsenNEWS.de (ed) (2021) *Anlagestrategie* [Online], Leipzig. Available at <https://www.boersennews.de/lexikon/begriff/anlagestrategie/113/> (Accessed 25 May 2021).
- Bortz, J. (2005) *Statistik für Human- und Sozialwissenschaftler*, 6th edn, Berlin, Heidelberg, New York, Springer.
- Breuer, W. and Breuer Claudia (2020) *Random-Walk-Hypothese* [Online]. Available at <https://wirtschaftslexikon.gabler.de/definition/random-walk-hypothese-44211/version-267527> (Accessed 4 December 2020).
- Camargo, A. and Smith, J. S. (2009) 'Image pattern classification for the identification of disease causing agents in plants', *Computers and Electronics in Agriculture*, vol. 66, no. 2, pp. 121–125.
- Chan, S. W.K. and Chong, M. W.C. (2017) 'Sentiment analysis in financial texts', *Decision Support Systems*, vol. 94, pp. 53–64.
- Cichosz, P. (2015) *Data mining algorithms: Explained using R / Paweł Cichosz, Department of Electronics and Information Technology, Warsaw University of Technology, Poland, Chichester, West Sussex, United Kingdom, Wiley*.
- Clark, P. and Boswell, R. (1991) 'Rule induction with CN2: Some recent improvements', in Kodratoff, Y. (ed) *Machine learning - EWSL - 91: [5th]European Working Session on Learning, Porto, Portugal, March 6-8, 1991 : proceedings / Y. Kodratoff (ed.)*, 482nd edn, Springer-Verlag Berlin, Heidelberg, pp. 151–163.
- Coenenberg, A. G., Salfeld, R. and Schultze, W. (2015) *Wertorientierte Unternehmensführung*, 3rd edn, Stuttgart, Schäffer-Poeschel.
- Dax-indices.com. (ed) (2020a) *DAX Digital | DAX® (TR) EUR*. [Online]. Available at <https://www.dax-indices.com/index-details?isin=DE0008469008> (Accessed 30 December 2020).

Dax-indices.com. (ed) (2020b) *DAX Digital | Zusammensetzung*. [Online]. Available at <https://www.dax-indices.com/de/web/dax-indices/zusammensetzung> (Accessed 29 December 2020).

Dax-indices.com. (ed) (2020c) *Factsheet DAX® Index*. [Online]. Available at https://www.dax-indices.com/document/Resources/Guides/Factsheet_DAX%20EUR_GR.pdf.

Dreger, C., Kosfeld, R. and Eckey, H.-F. (2013) *Ökonometrie: Grundlagen - Methoden - Beispiele*, 5th edn, Wiesbaden, Imprint: Springer Gabler.

Duda, R. O., Hart, P. E., Stork, D. G. and Duda, R. O. P. c. a. s. a. (2001) *Pattern classification*, 2nd edn, New York, Chichester, Wiley.

Fahrmeir, L., Heumann, C., Künstler, R., Pigeot, I. and Tutz, G. (2016) *Statistik: Der Weg zur Datenanalyse*, 8th edn, Springer Berlin Heidelberg.

Fama, E. F. (1970) 'Efficient Capital Markets: A Review of Theory and Empirical Work', *The Journal of Finance*, vol. 25, no. 2, pp. 383–417.

Field, A. P., Miles, J. and Field, Z. (2012) *Discovering statistics using R*, London, SAGE.

Finance.yahoo.com. (2020) *DAX PERFORMANCE-INDEX*. [Online]. Available at <https://finance.yahoo.com/quote/%5EGDAXI/history/> (Accessed 31 December 2020).

Finance.yahoo.com. (2021) *S&P 500 - Historical Data*. [Online]. Available at <https://finance.yahoo.com/quote/%5EGSPC/history/> (Accessed 5 January 2021).

Fox, J. and Weisberg, S. (2019) *R: An {R} Companion to Applied Regression* (Third Edition) [Computer program]. Available at <https://socialsciences.mcmaster.ca/jfox/Books/Companion/>.

Freed, M. and Lee, J. (2013) 'Application of Support Vector Machines to the Classification of Galaxy Morphologies', *International Conference on Computational and Information Sciences*, pp. 322–325.

Frost, I. (2017) *Einfache lineare Regression: Die Grundlage für komplexe Regressionsmodelle verstehen / Irasianty Frost*, Wiesbaden, Germany, Springer VS.

Fusion Media Limited (2020) *Investing.com* [Online]. Available at <https://www.investing.com/> (Accessed 4 April 2020).

Gehrke, M. (2019) *Angewandte empirische methoden in finance & accounting: Umsetzung mit r*, Boston, DeGruyter.

Gerlein, E. A., McGinnity, M., Belatreche, A. and Coleman, S. (2016) 'Evaluating machine learning classification for financial trading: An empirical approach', *Expert Systems with Applications*, vol. 54, pp. 193–207.

Gothein, W. (1995) *Evaluation von Anlagestrategien: Realisierung eines objektorientierten Simulators*, Wiesbaden, Deutscher Universitätsverlag; Imprint.

Han, J., Kamber, M. and Pei, J. (2012) *Data mining: Concepts and techniques / Jiawei Han, Micheline Kamber, Jian Pei*, 3rd edn, Waltham, MA, Morgan Kaufmann; [Oxford : Elsevier Science.

Hatekar, N. (2010) *Principles of Econometrics: An Introduction (Using R)*, B-42, Panchsheel Enclave, New Delhi 110 017 India, SAGE Publications India Pvt Ltd.

Hebbali, A. (2020) *R: olsrr: Tools for Building OLS Regression Models*. (R package version 0.5.3.) [Computer program]. Available at <https://cran.r-project.org/package=olsrr>.

IHS Markit (o.J.) *PMI TM by IHS Markit* [Online]. Available at <https://ihsmarkit.com/products/pmi.html> (Accessed 9 October 2020).

- John, G. H. and Langley, P. (1995) 'Estimating continuous distributions in Bayesian classifiers', in Besnard, P. and Hanks, S. (eds) *Uncertainty in artificial intelligence: Proceedings of the Eleventh conference on Uncertainty in artificial intelligence*, pp. 338–345.
- Kantardzic, M. (2020) *Data mining: Concepts, models, methods, and algorithms* / Mehmed Kantardzic, Hoboken, New Jersey, John Wiley & Sons.
- Karaca, Y. and Cattani, C. (2019) *Computational methods for data analysis*, Berlin Germany, Boston MA, De Gruyter.
- Kogan, P. (2009) *Marktanomalien bei IPOs: Frankfurter Freiverkehr und Londoner alternative investment market*, Hamburg, Igel-Verl.
- Koop, G. (2005) *Analysis of economic data*, 2nd edn, Chichester, Wiley.
- Kühn, S. and Kühn, M. (2019) *Alles über Fonds: Das Wichtigste zu ETF, Rendite, Kosten und Strategien : das Wissen aus dem Handbuch Geldanlage*, 2nd edn, Berlin, Stiftung Warentest.
- Kyereboah-Coleman, A. and Agyire-Tettey, K. F. (2008) 'Impact of macroeconomic indicators on stock market performance', *The Journal of Risk Finance*, vol. 9, no. 4, pp. 365–378.
- Lewis, D. D. (1998) 'Naive (Bayes) at forty: The independence assumption in information retrieval', in Nédellec, C. and Rouveirol, C. (eds) *Machine learning: ECML-98 : 10th European Conference on Machine Learning, Chemnitz, Germany, April 21-23 1998 : proceedings* / Claire Nédellec, Céline Rouveirol (eds.), Berlin, London, Springer, pp. 4–15.
- Mende A. (1990) *Erklärung des durchschnittlichen Kursniveaus an der New York Stock Exchange als Entscheidungshilfe für die Kapitalanlage in Aktien*, Freiburg, Albert-Ludwigs Universität.
- Mohammed, M., Khan, M. B. and Bashier, E. B. M. (2017) *Machine learning: Algorithms and applications*, Boca Raton, CRC Press Taylor & Francis Group.
- O'Brien, J. (2014) *Shaping knowledge: Complex social-spatial modelling for adaptive organizations*, Oxford, Chandos Publishing.
- Oppenländer, K.-H. (1996) *Konjunkturindikatoren: Fakten, Analysen, Verwendung*, 2nd edn, Oldenbourg Wissenschaftsverlag.
- Pilinkus, D. (2010) 'MACROECONOMIC INDICATORS AND THEIR IMPACT ON STOCK MARKET PERFORMANCE IN THE SHORT AND LONG RUN: THE CASE OF THE BALTIC STATES / MAKROEKONOMINIAI RODIKLIAI IR JŲ ĮTAKA AKCIJŲ RINKOS INDEKSUI TRUMPUOJU IR ILGUOJU LAIKOTARPIU: BALTIJOS ŠALIŲ ATVEJIS', *Technological and Economic Development of Economy*, vol. 16, no. 2, pp. 291–304.
- Plíhal, T. (2016) 'Stock Market Informational Efficiency in Germany: Granger Causality between DAX and Selected Macroeconomic Indicators', *Procedia - Social and Behavioral Sciences*, vol. 220, pp. 321–329.
- Qiu, D. (2015) *R: aTSA: Alternative Time Series Analysis*. (R package version 3.1.2.) [Computer program]. Available at <https://cran.r-project.org/package=aTSA>.
- Rastogi, R. and Shim, K. (2000) 'PUBLIC: A Decision Tree Classifier that Integrates Building and Pruning', *Data Mining and Knowledge Discovery*, vol. 4, no. 4, pp. 315–344.
- Schaberger, P. and Walker, A. (2020) *R: openxlsx: Read, Write and Edit xlsx Files*. (R package version 4.1.5.) [Computer program]. Available at <https://cran.r-project.org/package=openxlsx>.
- Schuster, T. (2017) *Statistik für Wirtschaftswissenschaftler*, 2nd edn, Springer Berlin Heidelberg.

- Shafaf, N. and Malek, H. (2019) 'Applications of Machine Learning Approaches in Emergency Medicine; a Review Article', *Archives of Academic Emergency Medicine*, vol. 7, no. 1.
- Shafer, J. C., Agrawal, R. and Mehta, M. (1996) 'SPRINT: A Scalable Parallel Classifier for Data Mining', in Vijayaraman, T. M. (ed) *Proceedings of the Twenty-Second International Conference on Very Large Data Bases, Mumbai (Bombay), India 3-6 September, 1996*, San Fransisco, Calif., Great Britain, Morgan Kaufmann Publishers, pp. 544–555.
- spglobal.com. (ed) (2021) *S&P 500® - S&P Dow Jones Indices*. [Online]. Available at <https://www.spglobal.com/spdji/en/indices/equity/sp-500/#overview> (Accessed 5 January 2021).
- Tan, P.-N., Steinbach, M., Karpatne, A. and Kumar, V. (2020) *Introduction to data mining*, NY, NY, Pearson.
- Tan, P.-N., Steinbach, M. and Kumar, V. (2006) *Introduction to data mining*, Harlow, Addison-Wesley.
- Theis, J. C. (2014) *Kommunikation zwischen Unternehmen und Kapitalmarkt: Eine theoretische und empirische Analyse von Informationsasymmetrien im Unternehmensumfeld*, Wiesbaden, Springer Gabler.
- TRADING ECONOMICS *Trading Economics* [Online]. Available at <https://tradingeconomics.com/> (Accessed 4 April 2020).
- Venitz, C. (2019) *Zum "value investing" aus Sicht der Bewertungstheorie*, Norderstedt, BoD – Books on Demand.
- Wildmann, L. (2016) *Wirtschaftspolitik: Module der Volkswirtschaftslehre*, Berlin, De Gruyter Oldenbourg.
- Winker, P. (2017) *Empirische Wirtschaftsforschung und Ökonometrie*, 4th edn, Berlin, Germany, Springer.
- Wolf, C. and Best, H. (2010) *Handbuch der sozialwissenschaftlichen Datenanalyse*, Wiesbaden, VS Verlag für Sozialwissenschaften.
- Wolf, J. (ed) (2009) *Methodik der empirischen Forschung*, 3rd edn, Wiesbaden, Gabler.
- Wooldridge, J. M. (2016) *Introductory econometrics: A modern approach*, 6th edn, South-Western, Cengage Learning.
- Xetra.com. (ed) (2020) *Deutsche Börse Group - DAX – Benchmark and Barometer for the German Economy*. [Online]. Available at <https://www.xetra.com/dbg-en/media/deutsche-boerse-spotlights/spotlight/DAX-benchmark-and-barometer-for-the-German-economy-139948> (Accessed 30 December 2020).
- Zantow, R. and Dinauer, J. (2011) *Finanzwirtschaft des Unternehmens: Die Grundlagen des modernen Finanzmanagements*, 3rd edn, München, Boston, Mass. [u.a.], Pearson Studium.
- Ziemba, W. T. (2012) *Calendar anomalies and arbitrage*, Hackensack N.J., World Scientific.

Abstract (Deutsch)

Diese Masterarbeit versucht unter Zuhilfenahme des Naive Bayes Klassifikators auf zwei unterschiedlichen Aktienmärkten, dem deutschen (DAX) sowie dem amerikanischen (S&P 500), die Rendite dieser Benchmarks im Betrachtungszeitraum 2004 bis 2019 zu übertreffen. Die Handelsentscheidungen werden basierend auf den a-posteriori Wahrscheinlichkeiten, abgeleitet auf Basis des Naive Bayes Klassifikators, getroffen. Den zweiten wesentlichen Eckpfeiler der entwickelten Strategie „Naive Bayes Investing“ bilden volkswirtschaftliche Indikatoren. Diese wurden aufgrund ihrer in der Literatur vielfach analysierten Interdependenzen mit dem Aktienmarkt als Inputfaktoren für den Naive Bayes Klassifikator und die Ableitung der Handelssignale ausgewählt. Aus einer Long List an volkswirtschaftlichen Indikatoren, die jeweils gesondert für den deutschen als auch amerikanischen Aktienmarkt abgeleitet wird, finden allerdings nur jene Indikatoren Eingang in Handelsentscheidungen, welche ausgewählte relevante statistische Kriterien erfüllen.

Hinsichtlich des Erfolgs der Strategie im Hinblick auf die Erzielung von Überrenditen im Vergleich zu den Benchmarks ergeben sich divergierende Ergebnisse. Während es beim DAX nicht gelang sich besser als der Markt zu stellen, zeigt sich beim S&P 500 ein gänzlich anderes Bild, da erhebliche Überrenditen erzielt werden konnten.

Ausgangspunkt zukünftiger Analysen könnte der Versuch sein, ein einheitliches Set an Inputvariablen zu finden, das eine Kombination aus makroökonomischen und technischen Indikatoren gleichermaßen berücksichtigt, und auf eine Vielzahl von Märkten anwendbar ist.

Abstract (English)

With the help of the Naive Bayes classifier, this master thesis attempts to generate a trading strategy that outperforms the returns of two selected stock markets, the German (DAX) and the American (S&P 500) stock market, over a period ranging from 2004 to 2019. Trading decisions are made based on the a-posteriori probabilities derived based on the Naive Bayes classifier.

Essential input parameters of the developed strategy "Naive Bayes Investing" are macroeconomic indicators. These were selected as input factors for the Naive Bayes Classifier and the derivation of trading signals due to their interdependencies with the stock market as widely discussed in literature. From a long list of economic indicators, which is derived separately for the German and American stock market, only those indicators that fulfil selected relevant statistical criteria are used in the derivation of trading decisions.

With regard to the success of the strategy in terms of achieving excess returns compared to the benchmarks, divergent results emerge. While the generated trading strategy did not succeed in outperforming the DAX (German market), it was effective for the S&P 500 (American market), as significant excess returns were achieved.

The starting point for future analyses could be the attempt to derive a uniform set of input variables that takes a combination of macroeconomic and technical indicators into account which are applicable to a large number of markets.