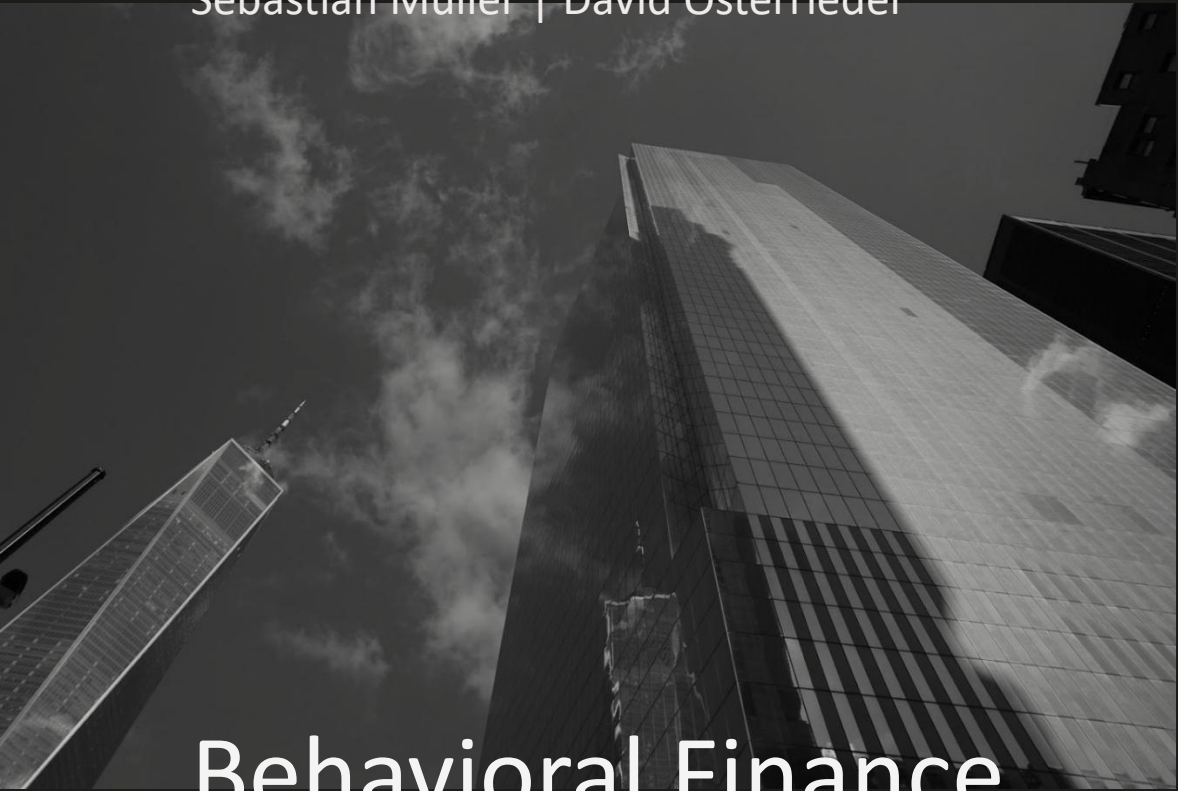


Künstliche Intelligenz, reale Auswirkungen – Wie generative KI Anleger, Märkte und das eigene Urteil verändert

Sebastian Müller | David Osterrieder



Behavioral Finance
Band 07

Künstliche Intelligenz, reale Auswirkungen – Wie generative KI Anleger, Märkte und das eigene Urteil verändert

Prof. Dr. Sebastian Müller und M.Sc. David Osterrieder

Disclaimer

Der vorliegende Band des Behavioral Finance e.V. dient ausschließlich der Information sowie dem besseren Verständnis technologischer und verhaltensökonomischer Entwicklungen an den Finanzmärkten. Die hierin getroffenen Aussagen, Analysen und Fallbeispiele zur Nutzung von Künstlicher Intelligenz und großen Sprachmodellen stellen in keiner Weise eine Anlageberatung, eine Kauf- oder Verkaufsempfehlung für Finanzinstrumente, eine Empfehlung für oder gegen einzelne Anbieter, Produkte oder Technologien dar.

Bitte bedenken Sie ausdrücklich, dass technologische Entwicklungen im Bereich der generativen KI einer hohen Dynamik und inhärenten Modellrisiken unterliegen. Genannte Zahlen, Marktbeispiele und Rechtsstände beziehen sich auf den Redaktionsstand, Juli 2026, und können zeitnah veralten. Die Betrachtungen in diesem Band lassen keine verlässlichen Rückschlüsse auf zukünftige Marktentwicklungen zu.

Aus Gründen der besseren Lesbarkeit verwenden wir in diesem Band das **generische Maskulinum** (etwa „der Anleger“). Gemeint sind damit **selbstverständlich Menschen aller Geschlechter**; die Formulierung ist rein sprachlich und ausdrücklich geschlechtsneutral zu verstehen.

© 2026
Sebastian Müller
David Osterrieder

Inhalt

1 Einleitung	1
2 Disruption der Welt durch generative KI	4
2.1 Privatanleger und Retail-Banken	5
2.2 Institutionelle Anleger und Anlageberatung	9
2.3 Finanzforschung	11
2.4 Agentic AI: Vom Chatbot zum handelnden Agenten	13
3 Risiken und psychologische Fallstricke der LLM-Interaktion	15
3.1 Determinanten des Modell-Outputs	16
3.2 Überzeugungskraft und Aversion	18
3.3 Korrektiv und Debiasing	19
4 Systemische und makroökonomische Implikationen	20
4.1 Markteffizienz und Trends im Asset Management	21
4.2 Struktur- und Jobwandel	23
4.3 Marktmanipulation und Betrugsszenarien	26
4.4 Herding 2.0: Das Risiko einer Modell-Monokultur	29
5 Zusammenfassung und Schlusswort	32
Literaturverzeichnis	35
Wir über uns	39
Change Log	41
Behavioral Finance Band 07	II



01

Einleitung

Ein Privatanleger tippt am Sonntagabend eine Frage in eine App: „Sollte ich jetzt in Technologieaktien investieren?“ und erhält Sekunden später eine flüssige, überzeugend formulierte Einschätzung. Wenige Bildschirme weiter wirbt ein gefälschtes, aber täuschend echtes Video eines bekannten Unternehmers für eine „garantiert profitable“ Krypto-Plattform. Beide Szenen haben dieselbe technologische Wurzel, und beide zeigen, warum generative Künstliche Intelligenz den Finanzsektor schneller verändert als nahezu jede andere Branche.

Finanzmärkte sind im Kern nichts anderes als Märkte für Informationen. Wer neue Daten schneller und präziser erfasst, bewertet und in Entscheidungen umsetzt, erzielt messbare Vorteile. Im Asset Management zeigt sich dies in Form von *Alpha*, einer systematischen Überrendite gegenüber dem Markt. Gleichzeitig entstehen auch jenseits der Anlageentscheidung Vorteile, etwa durch Effizienzgewinne bei standardisierten Prozessen wie KYC-Prüfungen, die Kosten senken und Abläufe beschleunigen können. Aus diesen Gründen war die Finanzindustrie schon immer extrem datengetrieben. Neu ist die Art der Daten, die nun maschinell verarbeitet werden kann.

Eine begriffliche Abgrenzung hilft, das Folgende einzuordnen. *Künstliche Intelligenz* (KI) ist der Oberbegriff für Systeme, die Aufgaben lösen, für die man normalerweise menschliche Intelligenz voraussetzt. *Maschinelles Lernen* (ML) ist der Teilbereich, in dem solche Systeme Muster aus Daten lernen, statt festprogrammiert zu werden. In der Finanzwelt geschah das traditionell vor allem quantitativ, also mit Charts, Bilanzkennzahlen, Analystenprognosen oder Ähnlichem.

Generative KI, insbesondere in Gestalt großer Sprachmodelle (englisch *Large Language Models*, kurz LLMs), bringt eine neue Qualität ins Spiel: Sie verarbeitet nicht mehr nur Zahlen, sondern auch Sprache und Bedeutung. Sie kann unstrukturierte Texte zusammenfassen, einordnen und eigenständig formulieren. Zu diesen Texten zählt auch Programmiercode; das Schreiben und Prüfen von Software ist inzwischen eines der größten Anwendungsfelder dieser Modelle.

Generative KI verschiebt die Grenze zwischen harten Zahlen und weichen Textinformationen, und genau dort entsteht das Spannungsfeld dieses Bandes.

Große Sprachmodelle sind im Kern verblüffend einfach gebaut: Sie sagen voraus, welcher nächste Wortbaustein am wahrscheinlichsten auf einen Text folgt, und lernen das, indem sie gewaltige Textmengen Satz für Satz „durchlesen“. Aus dieser stur wiederholten Übung entstehen Fähigkeiten, die niemand einzeln einprogrammiert hat: Zusammenfassen, Übersetzen, Finden, Einordnen, Antworten. Ein echtes „Verständnis“ im menschlichen Sinn ist das nicht; das Modell erzeugt statistisch plausible Sprache und kann deshalb auch überzeugend danebenliegen (für eine verständliche Einführung siehe Bank for International Settlements (2024)).

Der eigentliche Durchbruch liegt darin, *womit* diese Modelle arbeiten. Klassische Finanz-KI rechnet mit Zahlen: Kursen, Bilanzkennzahlen, Volumina. Der weitaus größte Teil der relevanten Finanzinformation liegt aber gar nicht als saubere Tabelle vor, sondern als Sprache: in Geschäftsberichten, Analystenkommentaren und -prognosen, Nachrichten, Notenbankprotokollen, makroökonomischen Berichten und politischen Meldungen oder im gesprochenen Wort einer Telefonkonferenz.

Schätzungen zufolge liegen 80 bis 90 Prozent aller Daten in unstrukturierter Form vor und waren damit maschinell kaum auswertbar (Harbert, 2021). Ihr Volumen wächst zudem rasant. Die weltweite Datenmenge dürfte laut IDC bis 2025 auf rund 175 Zettabyte gestiegen sein (Reinsel et al., 2018). 175 Zettabyte sind etwa 175 Billionen Gigabyte, umgerechnet also rund 22.000 Gigabyte für jeden Menschen auf der Erde. Dass sich dieser Schatz nun erschließen lässt, ist der eigentliche Hebel generativer KI im Finanzwesen.

Konkret können damit zahlreiche Aufgaben effizient bewältigt werden, die historisch mühsame Handarbeit erforderten:

- Geschäftsberichte und Pflichtmeldungen zusammenfassen, gezielt einzelne Fakten herausziehen und bewerten;
- das Sentiment tausender Nachrichten gleichzeitig bewerten;
- Telefonkonferenzen (*Earnings Calls*) nicht nur im Wortlaut, sondern bis in den Stimmklang auswerten;
- geldpolitische Äußerungen von Notenbanken analysieren und deren Ausrichtung als restriktiv („hawkish“) oder expansiv („dovish“) einordnen;
- ähnliche Unternehmen über Branchengrenzen hinweg finden und zu Netzwerken verknüpfen;
- Programmcode schreiben, Dokumente klassifizieren und übersetzen oder als Recherche-Assistent dienen;
- Routineprozesse automatisieren, etwa die Dokumentenprüfung zur Identitätsfeststellung (*Know Your Customer*, KYC) oder den Kundenservice.

Wie gewinnbringend LLMs eingesetzt werden können, zeigt eine aktuelle Studie: Lopez-Lira und Tang (2026) lassen ChatGPT Schlagzeilen danach bewerten, ob sie für eine Aktie gute oder schlechte Nachrichten sind und finden so einen positiven Zusammenhang mit folgenden Kursbewegungen.

In diesem Band untersuchen wir, wie generative KI den Finanzbereich verändert. Unser Schwerpunkt liegt dabei auch auf der Behavioral Finance, also der verhaltenswissenschaftlich fundierten Finanzforschung, ohne dabei die systemischen Implikationen auszublenden. Unter anderem interessiert uns die Frage, was die Technik mit den Menschen macht, die sie nutzen. Dafür zeigen wir in Kapitel 2, wie generative KI die Arbeit der verschiedenen Marktakteure (vom Privatanleger über den Anlageberater bis zum Finanzmarktforscher) bereits heute verändert. Kapitel 3 untersucht die psychologischen Fallstricke der Mensch-Maschine-Interaktion. Kapitel 4 weitet den Blick auf die systemische Ebene: Von der Effizienz der Märkte und Trends auf dem Arbeitsmarkt über KI-gestützten Betrug bis zu der Frage, was passiert, wenn viele dieselben Modelle nutzen. Kapitel 5 fasst zusammen und gibt praktische Leitplanken.

Unser Ziel ist es, generative KI im Finanzbereich nüchtern einzuordnen und somit ihre Chancen greifbar zu machen, ihre Fallstricke offen zu benennen und eine praktische Orientierung für den eigenen Umgang mit diesen Werkzeugen zu geben.



Disruption der Welt durch generative KI

Dieselbe Technologie wirkt in jedem Bereich des Finanzsektors individuell. Ein Sprachmodell, das einem Privatanleger den Einstieg erleichtert, erfüllt im institutionellen Fondsmanagement oder in der Forschung völlig andere Zwecke. Um die realen Konsequenzen greifbar zu machen, betrachtet dieses Kapitel vier Perspektiven und schließt mit einem Blick auf den aktuellsten Entwicklungsschritt, die sogenannten KI-Agenten. Dabei zählt nicht nur, was die Technik kann, sondern auch, wie sie das Verhalten der Beteiligten verändert.

2.1 Privatanleger und Retail-Banken

Für Privatanleger senkt generative KI die Schwelle zu Finanzwissen und Marktanalyse drastisch. Sprachmodelle können etwa Geschäftsberichte, Quartalszahlen oder Ad-hoc-Mitteilungen zusammenfassen, zentrale Kennzahlen erklären und gezielt auf Rückfragen des Nutzers eingehen. Was lange vermögenden Kunden in der Vermögensverwaltung vorbehalten war, nämlich eine verständliche und auf die eigene Situation zugeschnittene finanzielle Einordnung, liefert heute ein Chatfenster in Sekunden. Auch die digitalen Anlagehelfer, die *Robo-Advisor*, verändern sich: Wo früher starre Online-Fragebögen und feste Allokationsregeln standen, treten zunehmend dialogbasierte Oberflächen, die sich wie ein Messenger bedienen lassen. Selbst etablierte Dienste öffnen sich dafür: In Google Finance etwa können Nutzer seit Ende 2025 mithilfe des Sprachmodells Gemini in natürlicher Sprache Fragen zu einem Unternehmen stellen, während sie dessen Kennzahlen und Charts vor sich haben (Google, 2025). Einzuordnen ist das jedoch mit Vorsicht. Die etablierten Robo-Advisor nutzen KI bislang vor allem im Hintergrund, etwa zur Optimierung der Portfolios, während eine wirklich dialogbasierte, regulierte KI-Beratung im Privatkundengeschäft noch am Anfang steht. Fest steht aber, dass der einfachere Zugang, gepaart mit spielerischen Elementen (*Gamification*), die Eintrittsbarrieren senkt, insbesondere für jüngere und technisch versiertere Anleger.

Besonders wertvoll ist diese Zugänglichkeit bei komplexen Finanzprodukten. Wer früher vor einem Wertpapierprospekt, einem mit Fachbegriffen überladenen Termsheet oder einem strukturierten Zertifikat kapitulierte, kann sich heute Schritt für Schritt erklären lassen, wie ein Hebelprodukt funktioniert, wo die Kostenfallen eines aktiv gemanagten Fonds liegen oder was eine Vorfälligkeitsentschädigung bei der Baufinanzierung bedeutet. Sprachmodelle beantworten Rückfragen geduldig, in einfacher Sprache und so oft, wie der Nutzer es braucht, ohne dass dieser sich vor einem Berater bloßgestellt fühlt. Gerade darin liegt ein unterschätzter Wert: Wer aus Unsicherheit gar nicht erst investiert und sein Geld dauerhaft auf dem Girokonto liegen lässt, verliert real an Kaufkraft und verzichtet über die Jahre auf erhebliche Renditechancen. Generative KI kann diese Verständnishürde senken und damit dazu beitragen, die Opportunitätskosten des Nichtstuns zu reduzieren.

Damit liegt ein besonderer Gewinn generativer KI in der *Financial Literacy*, also im Finanzwissen. Sprachmodelle können Verständnishürden senken, indem sie nicht nur Informationen bereitstellen, sondern diese auf Nachfrage, in einfacher Sprache und entlang des konkreten Problems erklären. Finanzprodukte werden dadurch nicht automatisch besser, aber sie werden zugänglicher: Anleger können Kosten, Risiken und Anreize besser nachvollziehen und dadurch informiertere Entscheidungen treffen.

Wirtschaftlich ist das durchaus relevant. Wer einfache Kosten- und Risikostrukturen versteht, erkennt eher, wann Produkte unnötig komplex oder teuer sind. Aktiv gemanagte Fonds schlagen ihren Vergleichsindex über lange Zeiträume nur selten; laut der SPIVA-Scorecard von S&P liegt über fünfzehn Jahre die große Mehrheit der aktiven US-Aktienfonds hinter dem Markt (S&P Dow Jones Indices, 2025), und Fama und French (2010) führen dies überwiegend auf Kosten und Zufall statt auf echtes Können zurück. Passend dazu zeigt eine Studie, an der einer von uns beteiligt ist, dass bei der Fondswahl *Einfachheit* gewinnt, und zwar unabhängig vom Finanzwissen der Anleger (Gong & Müller, 2026).

Genau an dieser Stelle wird generative KI für Privatanleger interessant: Sie kann einfache, robuste Grundprinzipien verständlich machen und auf konkrete Lebenssituationen herunterbrechen. Ein kurzer Praxistest mit aktuellen Sprachmodellen zeigt, wie niedrig die Einstiegshürde inzwischen ist. Auf eine einfache Frage zur privaten Altersvorsorge geben verschiedene Modelle keine spekulativen Aktientipps, sondern verweisen auf Grundprinzipien wie einen langen Anlagehorizont, breite Streuung, niedrige Kosten und einen ETF-Sparplan als naheliegende Basis.

Beispiel: Wie die führenden KI-Modelle auf eine Anlagefrage antworten

Nutzerfrage: „Ich bin 30, verdiene durchschnittlich 2700€ netto im Monat und möchte davon 200€ im Monat für die Altersvorsorge anlegen. Was rätst du mir?“

Antwort Google Gemini (gekürzt): „Wichtiger Hinweis vorab: Ich bin eine KI und kein zertifizierter Finanzberater. Dies sind allgemeine Informationen und keine individuelle Anlageberatung. Hier ist ein bewährtes Basiskonzept für deine Situation: Für einen Anlagehorizont von über 30 Jahren empfehlen Verbraucherzentralen und Finanzexperten in der Regel passive Indexfonds (ETFs). Sie sind günstig, pflegeleicht und sehr transparent. Anstatt zu versuchen, einzelne „Gewinner-Aktien“ zu finden, kaufst du mit einem Welt-ETF (wie dem MSCI World oder dem FTSE All-World) einen kleinen Anteil von tausenden Unternehmen weltweit. Die historische Durchschnittsrendite des weltweiten Aktienmarktes lag über lange Zeiträume (trotz aller Krisen) bei etwa 7% bis 8% pro Jahr nach Inflation. [...]“

Antwort OpenAI GPT-5.5 Pro (gekürzt): „Bei 30 Jahren und einem Anlagehorizont von ca. 35–37 Jahren ist für die Altersvorsorge meist ein einfacher, breit gestreuter ETF-Sparplan die beste Basis — vorausgesetzt, du hast keinen teuren Konsumkredit und einen Notgroschen. [...] Keine individuelle Anlage- oder Steuerberatung, aber als einfache, robuste Strategie wäre das für deine Situation ein sehr vernünftiger Start. [...]“

Unsere Einordnung: Zweierlei fällt auf. Erstens ziehen die Modelle eine klare Grenze zur individuellen Anlageberatung und ordnen ihre Antworten als allgemeine Information ein. Das ist wichtig, weil eine konkrete Anlageempfehlung von der persönlichen Risikotragfähigkeit, bestehenden Schulden, Liquiditätsreserven und steuerlichen Situation abhängt. Zweitens ist der inhaltliche Rat bemerkenswert solide. Beide Modelle empfehlen keine Einzeltitel, keine aktiven Fonds oder komplexen Produkte, sondern eine einfache, breit gestreute und kostengünstige Lösung. Genau das entspricht für die meisten Privatanleger dem, was die empirische Forschung nahelegt. Lediglich einzelne Zahlenangaben müssen geprüft werden, denn die von Gemini genannte globale Rendite von 7 bis 8 Prozent nach Inflation ist etwas optimistisch.

Gerade darin liegt der zentrale Nutzen generativer KI für Privatanleger: Sie macht nicht nur Finanzwissen leichter zugänglich, sondern kann auch dabei helfen, gute Standardprinzipien der Finanzforschung in konkrete, verständliche Handlungsoptionen zu übersetzen.

Der Einsatz generativer KI ist aber nicht auf Erklärung und Orientierung beschränkt. Dieselben Modelle, die einen Geschäftsbericht verständlich zusammenfassen oder die Logik eines ETF-Sparplans erklären, können auch Nachrichten, Analystenkommentare, Social-Media-Beiträge oder Unternehmensmitteilungen auswerten und daraus vermeintliche Stimmungssignale für einzelne Aktien ableiten. Für Privatanleger ist das besonders verführerisch: Was früher Datenzugang, Programmierkenntnisse und institutionelle Infrastruktur erforderte, lässt sich heute scheinbar über eine einfache Eingabeaufforderung abrufen.

Dass solche Stimmungssignale tatsächlich existieren, deutet bereits die in der Einleitung erwähnte Studie von Lopez-Lira und Tang (2026) an. Für den Privatanleger ist aber auch eine nüchterne Einordnung entscheidend. Die teils spektakulären Renditen dieser Studie gelten vor Handelskosten und sind bei kleinen, wenig liquiden Aktien am stärksten ausgeprägt; sie zeigen sich aber auch bei größeren, liquideren Titeln. Hinzu kommt, dass mit zunehmender Verbreitung der Sprachmodelle der Vorteil im Zeitverlauf deutlich schmilzt. Die risikoadjustierte Performance (*Sharpe-Ratio*) sinkt über wenige Jahre von sehr hohen auf deutlich niedrigere Werte. Das ist ein erster Hinweis auf den *Alpha Decay*, der uns in Kapitel 4 noch beschäftigt.

Für Privatanleger folgt daraus eine wichtige Unterscheidung. Dass KI-Modelle in bestimmten Situationen tatsächlich nützliche Signale erzeugen können, macht sie noch nicht zu einer robusten Anlagestrategie für den Alltag. Gerade weil einzelne Antworten plausibel und manchmal empirisch hilfreich sein können, liegt die Gefahr darin, ihre Verlässlichkeit zu überschätzen.

Die sprachliche Sicherheit eines Modells lässt oft weniger Zweifel erkennen, als bei Anlageentscheidungen eigentlich angebracht wären. Dadurch können Anleger nicht nur die Qualität einzelner KI-Hinweise überschätzen, sondern auch ihre eigene Fähigkeit, diese Signale richtig einzuordnen und gewinnbringend zu nutzen. Je einfacher solche Hinweise abrufbar sind und je schneller der Weg von der Information zur Order wird, desto größer wird die Gefahr von *Overconfidence* und übermäßigem Handeln (*Overtrading*).

Die niedrige Eintrittsbarriere ist Chance und Risiko zugleich: Sie demokratisiert den Zugang - und kann zu häufigerem, impulsiverem Handeln verleiten.

Dass häufiges Handeln der Rendite schadet, ist seit Langem belegt. Barber und Odean (2000) zeigen für über 60.000 Privatanleger, dass die aktivsten Händler am deutlichsten hinter dem Markt zurückbleiben; ein erheblicher Teil dieses Renditenachteils ließ sich damals durch Handelskosten erklären. Neuere Evidenz zur Robinhood-Generation macht jedoch deutlich, dass das Problem auch bei deutlich niedrigeren Transaktionskosten fortbesteht. Aufmerksamkeitsgetriebenes Handeln, wie es Apps mit Ranglisten und Push-Nachrichten fördern, führt zu Herdenkäufen und im Schnitt schlechteren Renditen (Barber et al., 2022).

Der eigentliche Kern bleibt damit ein Verhaltensproblem. Wer ständig neue Signale sieht, ständig neue Erklärungen erhält und ständig neue Handlungsoptionen angezeigt bekommt, handelt leichter überaktiv und aus Selbstüberschätzung. Hinzu kommt ein steuerlicher Nachteil: Jedes Realisieren von Gewinnen kann Steuern auslösen, die bei ruhigerem Halten noch lange gestundet blieben und weiter mitarbeiten könnten. Eine KI, die rund um die Uhr neue Ideen liefert und jeden Impuls sofort handelbar erscheinen lässt, kann genau dieses Muster verstärken.

Dass solche Werkzeuge im Privatanlegersegment bereits ankommen, zeigt eine Eigenbefragung des Brokers eToro unter rund 11.000 Anlegern im Oktober 2025. Dort gaben 19 Prozent der Befragten an, KI-Werkzeuge für ihre Anlageentscheidungen zu nutzen (eToro, 2025). Mit der zunehmenden Integration solcher Funktionen in Broker-Apps, Suchmaschinen und Finanzportale dürfte dieser Anteil weiter steigen.

Die Bilanz wäre jedoch unvollständig, wenn KI nur als Verstärker impulsiven Handelns beschrieben würde. Zur *Financial Literacy* gehört auch, Warnsignale zu erkennen und fragwürdige Anbieter einordnen zu können. Wer unsicher ist, ob ein Anbieter seriös ist, kann ein Sprachmodell nutzen, um erste Prüfungen zu strukturieren: Existiert das Unternehmen überhaupt, ist es zugelassen, gibt es behördliche Warnungen oder auffällige Erfahrungsberichte?

Solche Prüfungen ersetzen keine eigene Sorgfalt und keine offiziellen Register, können aber eine erste, schnelle Möglichkeit sein, bekanntere Betrugsversuche früh zu erkennen. Wie das praktisch gelingt und wo die Grenzen liegen, vertiefen wir in Kapitel 4.3.

Zuletzt muss erwähnt werden, dass sich trotz sinkender Hürden zugleich eine Schere öffnet. Wer sprachlich sicher, technisch versiert und zahlungswillig für Premium-Modelle ist, profitiert von der Entwicklung überproportional. Ältere Anleger, Menschen mit niedrigerem Bildungsstand oder in Regionen mit schlechterem Internetzugang drohen komparativ zurückzufallen. Generative KI demokratisiert Wissen – aber nur für jene, die die Eintrittsschwelle überhaupt überwinden.

2.2 Institutionelle Anleger und Anlageberatung

In der professionellen Beratung verschiebt sich der Fokus von der Demokratisierung hin zu Effizienz. KI entlastet die Beschäftigten hier zunächst von Routinearbeit. Ein prägnantes Beispiel ist Morgan Stanley: Die Bank führte 2023 in Kooperation mit OpenAI einen internen Wissensassistenten ein, der den Beratern Antworten aus rund 100.000 internen Research-Dokumenten liefert, ohne auf das offene Internet zuzugreifen (Morgan Stanley, 2023). Ein Jahr später folgte ein Werkzeug, das mit Einwilligung des Kunden Gespräche protokolliert und Zusammenfassungen erstellt. Der Berater gewinnt Zeit für das Wesentliche: die persönliche Führung des Mandanten. Ähnliche Assistenten fassen inzwischen Geschäftsberichte und Analystenkonferenzen zusammen, liefern erste Entwürfe für Memos oder Pitch-Unterlagen, durchsuchen Verträge und unterstützen die Geldwäsche- und Compliance-Prüfung.

Auch in der hauseigenen Technik greift die Veränderung tief. Goldman Sachs etwa stellte 2025 einen KI-Assistenten firmenweit allen Mitarbeitenden bereit (Fox Business, 2025) und erprobt autonom arbeitende Software-Agenten, um etwa veralteten Programmcode auf neue Versionsstände zu heben (CNBC, 2025). Solche Effizienzgewinne im Hintergrund setzen Gelder frei und verändern mittelfristig auch die Anforderungen an die Beschäftigten.

Neben der Effizienz rückt zunehmend die Schadensvermeidung in den Blick, also der Wert, der erhalten bleibt, weil ein Problem früh erkannt wird. Sprachmodelle werten rund um die Uhr Nachrichten, Geschäftsberichte und Behördenmitteilungen aus und können als Frühwarnsystem dienen, etwa wenn sich bei einem wichtigen Zulieferer rechtliche Probleme, Streiks oder Lieferengpässe andeuten.

So lassen sich Klumpen- und Lieferkettenrisiken eines Portfolios sichtbar machen, oft bevor sie in den Kursen stehen. Im stark regulierten Tagesgeschäft können Modelle zudem prüfintensive Massenaufgaben wie etwa die Identitätsprüfung von Neukunden (*Know Your Customer*), das Sichten von Transaktionen auf Geldwäsche und Betrug oder die Aufbereitung von Kreditentscheidungen (*Credit Scoring*) übernehmen. Der große Hebel liegt dabei nicht nur in höherer Geschwindigkeit, sondern auch in Skalierbarkeit und gleichbleibender Aufmerksamkeit. Gerade standardisierte Prüfungen ermüden menschliche Bearbeiter, obwohl kleine Auffälligkeiten entscheidend sein können. Modelle zeigen keine Ermüdungserscheinungen und können dieselbe Prüflogik über viele Fälle hinweg konsistent anwenden. Hinzu kommt, dass viele dieser Aufgaben nicht sequenziell aufeinander aufbauen: Tausende KYC-Prüfungen, Transaktionschecks oder Dokumentensichtungen können nun parallel bearbeitet werden. Wo früher zusätzliche Fallzahlen zusätzliche Teams erforderten, lassen sich heute mehrere Modelle oder Agenten, mehr dazu später, gleichzeitig einsetzen.

Je näher KI-Systeme an Kundendaten, internen Research-Dokumenten oder regulatorisch relevanten Entscheidungen arbeiten, desto wichtiger werden jedoch geschlossene, kontrollierte Umgebungen. Ein Grund dafür ist die *Vertraulichkeit*. Wer Kundendaten oder unveröffentlichte Analysen in ein öffentliches Sprachmodell eingibt, gibt sie unter Umständen aus der Hand. Für Privatanleger gilt im Kleinen dasselbe: Sensible Finanzdaten gehören nicht ungeprüft in ein fremdes Chatfenster.

Hinzu kommt eine zweite wichtige Grenze. So komfortabel die maschinellen Zuarbeiten sind, ersetzen sie nicht die analytische Verantwortung des Menschen. Wer den überzeugend formulierten Auswertungen seines KI-Assistenten zunehmend vertraut, delegiert schleichend einen Teil der eigenen Prüfung. Genau in diesem stillen Vertrauen liegt ein verhaltensökonomisches Risiko, dem sich Kapitel 3 widmet. Dass dieselben Effizienzgewinne mittelfristig auch Tätigkeiten und ganze Stellen verändern oder verdrängen können, ist die andere Seite dieser Entwicklung, die wir in Kapitel 4 aufgreifen.

2.3 Finanzforschung

Auch in der akademischen Forschung löst generative KI einen methodischen Sprung aus. Empirische Kapitalmarktstudien stützten sich lange fast ausschließlich auf harte Zahlenreihen. Die systematische Auswertung weicher, narrativer Informationen scheiterte an der schieren Masse an Text. Sprachmodelle schließen diese Lücke.

Ganz neu ist die Idee, Texte auszuwerten, allerdings nicht. Schon vor den Sprachmodellen zählte man Wörter mit festen Listen aus, der *wörterbuchbasierte* Ansatz: Ein vorab definiertes Lexikon legt fest, welche Begriffe als positiv oder negativ gelten. Loughran und McDonald (2011) zeigen allerdings, wie sehr es auf die richtige Wortliste ankommt. Ein eigens für Finanzberichte gebautes Lexikon schneidet deutlich besser ab als ein allgemeinsprachliches, in dem Wörter wie „liability“ (Verbindlichkeit) oder „tax“ (Steuer) fälschlich als negativ gewertet werden. Weiterhin bleibt der Haken solcher Verfahren, dass sie nur Wörter zählen und weder Zusammenhang noch Verneinung oder Ironie verstehen. Der Satz „Von einer Erholung der Gewinne kann keine Rede sein.“ etwa ist für einen Menschen klar negativ, enthält aber kein einzelnes eindeutig negatives Wort, an dem eine starre Liste das festmachen könnte. Genau hier setzen generative Sprachmodelle an, die Bedeutung im Kontext erfassen.

Wie produktiv das sein kann, zeigt eine Studie, an der einer von uns beteiligt ist: Breitung und Müller (2025). Dort erzeugt ein Sprachmodell zunächst für jedes Jahr eine *historisch akkurate*, also nur mit dem Wissen, das zum jeweiligen Zeitpunkt verfügbar war, Geschäftsbeschreibung für über 63.000 Unternehmen und konstruiert daraus Ähnlichkeits- und Verflechtungsnetzwerke. Mit ihnen lassen sich anschließend unter anderem Kurszusammenhänge zwischen verbundenen Firmen oder potenzielle Übernahmeziele aufspüren.

Methodisch heikel ist bei solchen Anwendungen jedoch ein möglicher *Lookahead-Bias* (Rückschaufehler). Das Problem lässt sich am besten anhand eines Beispiels veranschaulichen. Soll ein Modell aus Sicht des Jahres 2018 beurteilen, ob ein Unternehmen attraktiv bewertet ist, darf es nicht wissen, dass dessen Aktie 2019 stark gefallen ist oder das Unternehmen 2020 übernommen wurde. Moderne Sprachmodelle wurden aber häufig auf Texten trainiert, die weit über den historischen Untersuchungszeitpunkt hinausreichen. Lässt man sie vergangene Lagen bewerten, können spätere Entwicklungen ungewollt in die Antwort einfließen. Eine Strategie wirkt dann im Rückblick treffsicherer, als sie es in Echtzeit gewesen wäre.

Besonders schwierig ist, dass Sprachmodelle in dieser Hinsicht eine *Black Box* bleiben. Selbst wenn man sie ausdrücklich anweist, späteres Wissen nicht zu verwenden und nur aus der Perspektive des damaligen Zeitpunkts zu antworten, lässt sich nicht zuverlässig kontrollieren, ob dieses Wissen intern doch eine Rolle spielt. Deshalb reicht es in der Forschung nicht aus, ein Modell einfach per Anweisung „so tun“ zu lassen, als sei es 2018, sondern zusätzliche Schutzmaßnahmen sind nötig.

Eine erste Grundregel ist, den Internetzugang des Modells für solche Auswertungen abzuschalten. Sonst kann es sich spätere Informationen nachträglich beschaffen, selbst wenn sein internes Trainingswissen begrenzt ist. Darüber hinaus gibt es mehrere Strategien, um Lookahead-Bias zu reduzieren. Ein Ansatz ist *Masking*. Dabei werden Firmennamen, konkrete Angaben oder andere eindeutige Hinweise entfernt oder mit generischen Platzhaltern ersetzt, damit das Modell später bekannt gewordene Ereignisse weniger leicht abrufen kann. Ein zweiter Ansatz nutzt den *Knowledge Cutoff*, also den Stichtag, bis zu dem ein Modell Informationen in seinem Training gesehen haben kann. Forschende können ein Modell oder einen Anbieter gezielt so auswählen, dass dieser Stichtag vor dem Zeitpunkt liegt, aus dessen Perspektive die historische Lage bewertet werden soll. Soll ein Modell etwa eine Lage aus dem Jahr 2018 bewerten, darf sein Trainingswissen keine Informationen aus späteren Jahren enthalten. In der Praxis ist dieser Ansatz jedoch begrenzt, denn gerade die leistungsfähigsten Modelle haben meist neuere Wissensstände, und selbst ältere verfügbare Modelle reichen häufig nicht weit genug zurück, um Analysen vor den frühen 2020er-Jahren ohne Risiko von Lookahead-Bias durchzuführen.

Einen systematischeren Weg bieten deshalb *Point-in-time*-Modelle. Ihr Vorteil liegt nicht nur in einem früheren Wissensstand, sondern in der kontrollierten Konstruktion einer ganzen Modellreihe. Idealerweise werden Modelle mit gleicher Architektur, aber unterschiedlichen zeitlichen Stichtagen trainiert und versioniert, etwa monatlich oder jährlich. Dadurch lassen sich historische Analysen über verschiedene Zeitpunkte hinweg besser vergleichen. Zudem können solche Modelle lokal gespeichert und später erneut ausgeführt werden; der Zugriff hängt also nicht davon ab, ob ein kommerzieller Anbieter ein bestimmtes altes Modell weiterhin bereitstellt oder verändert. Kelly et al. (2026) verfolgen diesen Ansatz mit zeitgestempelten Sprachmodellen und rücken historische Analysen damit näher an die tatsächliche Informationslage eines damaligen Marktteilnehmers heran.

2.4 Agentic AI: Vom Chatbot zum handelnden Agenten

Dieser Entwicklungsschritt geht über das antwortende Sprachmodell hinaus. Bei der *Agentic AI* werden mehrere Schritte verkettet: Ein KI-Agent kann eigenständig recherchieren, Werkzeuge bedienen und Aktionen ausführen. In den Großbanken ist diese Entwicklung angekommen; interne KI-Plattformen werden inzwischen breit ausgerollt. Auch im Privatkundengeschäft wird es konkret: Im Mai 2026 öffnete der Broker Robinhood seine Plattform in einer Beta-Version erstmals für KI-Agenten (Robinhood, 2026). Das Unternehmen zählt nach eigenen Angaben rund 27 Millionen Kundenkonten mit Guthaben. Nutzer können einen externen Agenten (etwa auf Basis gängiger Sprachmodelle) über eine standardisierte Schnittstelle (das *Model Context Protocol*, kurz MCP) anbinden, woraufhin dieser eigenständig Aktien handelt. Aus Sicherheitsgründen geschieht das zunächst nur in einem abgetrennten, separat befüllten Konto, jede Order löst eine Benachrichtigung aus, und der Zugang lässt sich jederzeit kappen. Das Risiko verbleibt ausdrücklich beim Nutzer, und Robinhood prüft oder überwacht die fremden Agenten nicht. Noch ist das ein vorsichtiger Anfang, beschränkt auf Aktien und ausdrücklich als Beta gekennzeichnet.

Wie sich solche Agenten an einem Markt verhalten, lässt sich aber auch erforschen, ohne echtes Geld zu riskieren. Lopez-Lira (2025) lässt in einer quelloffenen Simulation mehrere Sprachmodelle als konkurrierende Händler gegeneinander antreten. Dabei agieren die einen als Value-Investoren, die anderen als Momentum-Trader oder als Market-Maker. Zwei Befunde sind bemerkenswert: Erstens halten die Modelle ihre zugewiesene Strategie über viele Handelsrunden hinweg konsequent durch, agieren also als belastbare, eigenständige Marktteilnehmer. Zweitens bringt der künstliche Markt von selbst realistische Muster hervor, etwa Preisfindung, spekulative Blasen und träge Reaktionen auf Nachrichten. Solche Laborumgebungen erlauben es, Finanztheorien zu testen, für die es keine geschlossene Lösung gibt. Außerdem sind sie deutlich günstiger als Experimente mit menschlichen Probanden.

Ein weiterer Vorzug betrifft die Anreize. Während ein menschlicher Proband je nach Prämie unterschiedlich aufmerksam mitarbeitet oder die Aufgabe missversteht, antwortet ein KI-Agent gleichbleibend, ob nun fünf oder fünfhundert Euro Belohnung im Spiel sind, sodass verzerrte Anreize die Ergebnisse weniger trüben. Zugleich legen sie eine Kehrseite offen, auf die wir in Kapitel 4.4 zurückkommen: Ähnlich formulierte Anweisungen können die Agenten zu gleichgerichtetem Verhalten verleiten.

Autonomes Handeln durch KI ist real, aber noch die Ausnahme: Aufsichtsbehörden stufen bislang nur einen kleinen Teil der Anwendungen als vollständig autonom ein.

Hier ist jedoch ein nüchterner Realitätscheck angebracht. Nach einer gemeinsamen Erhebung der Bank of England und Finanzaufsicht FCA zum britischen Finanzsektor waren zuletzt nur etwa 2 Prozent der KI-Anwendungsfälle vollständig autonom (Bank of England & Financial Conduct Authority, 2024). Der ganz überwiegende Teil arbeitet folglich weiterhin unter menschlicher Aufsicht. Mit zunehmender Autonomie verändert sich jedoch die zentrale Herausforderung: Es geht dann nicht mehr nur darum, ob eine Maschine zutreffende Einschätzungen liefert, sondern auch darum, welche Handlungen sie selbstständig ausführt, wie diese kontrolliert werden können und wer für ihre Folgen verantwortlich ist.



03

Risiken und psychologische Fallstricke der LLM- Interaktion

Sobald Finanzakteure Analyse- und Rechercheaufgaben an Sprachmodelle delegieren, ändert sich die Psychologie der Entscheidungsfindung. Wir kommunizieren nicht mehr mit einer stummen Datenbank, sondern mit einem System, das flüssig und in sich schlüssig argumentiert. Diese fast menschliche Interaktion weckt bekannte kognitive Verzerrungen. Dieses Kapitel zeigt, welche das sind und wie ein wenig technisches Verständnis helfen kann, potenziell folgenschwere Fehlentscheidungen zu vermeiden.

3.1 Determinanten des Modell-Outputs

Die Qualität einer KI-Antwort hängt von vielen Details ab, im Kern aber von zwei Hebeln: davon, *wie das Modell entstanden ist* (mit welchen Trainingsdaten und mit welchem Feinschliff), und davon, *wie es eingesetzt wird* (welches Modell man wählt, wie man fragt, welche Hintergrundinformationen man mitgibt und auf welche Werkzeuge und Datenquellen es zugreifen darf).

Zunächst gilt das alte Prinzip der Datenverarbeitung: „Garbage In, Garbage Out“. Ein Sprachmodell besitzt kein echtes Verständnis von Märkten, sondern es berechnet Wahrscheinlichkeiten von Wortfolgen auf Basis seiner Trainingstexte. Wurde es überwiegend mit Daten aus einer langen Aufschwungphase trainiert, kann es zum Beispiel zu optimistischeren Einschätzungen neigen.

Der zweite Faktor ist der *Prompt*, also die Eingabe oder Anweisung an die KI. Hier greift ein klassisches Konzept der Verhaltensökonomie: der *Framing-Effekt*. Tversky und Kahneman (1981) zeigen, dass schon die Formulierung einer Frage die Entscheidung beeinflusst. Auf das *Prompt-Engineering* (das gezielte Gestalten der Eingaben an die KI) übertragen heißt das: Wer fragt „Welche Risiken sprechen gegen dieses Investment?“, erhält eine ganz andere Antwort als bei „Welche Chancen bietet dieses Investment?“. Das Modell spiegelt die Prämisse der Frage. Ohne neutrale Fragestellung kann die KI so zum Verstärker des eigenen *Bestätigungsfehlers* (englisch *Confirmation Bias*) werden. Sie liefert dann genau die Antwort, die die vorgefasste Meinung stützt, auch wenn diese Meinung nur indirekt aus der Formulierung der Frage abgeleitet wird.

Über die Frage hinaus entscheidet zunehmend, *worauf* ein Modell zugreifen darf. Das ist besonders relevant im Kontext des Knowledge Cutoffs, den wir bereits in Kapitel 2.3 genauer beleuchtet haben. Informationen nach diesem Wissensstichtag sind dem Modell nicht automatisch bekannt. Sie können nur dann in die Antwort einfließen, wenn sie der Nutzer bereitstellt oder wenn das Modell Zugriff auf aktuelle externe Quellen hat. Gerade im Finanzbereich ist das entscheidend: Eine Antwort kann sprachlich aktuell wirken, obwohl sie auf einem veralteten Informationsstand beruht.

Erhält ein Modell dagegen Zugriff auf Live-Daten, Datenbanken oder andere digitale Werkzeuge, verändert sich seine Rolle. Es beantwortet dann nicht mehr nur Fragen auf Basis seines Trainingswissens, sondern kann aktuelle Kurse abrufen, im Internet recherchieren, Dokumente durchsuchen oder Berechnungen ausführen. Neue standardisierte Schnittstellen, allen voran das bereits erwähnte MCP, erleichtern es, Sprachmodelle mit Fachprogrammen, internen Datenquellen und externen Informationsdiensten zu verbinden. Das macht die Systeme leistungsfähiger, verschiebt aber auch Verantwortung zum Nutzer oder zur Organisation, die diese Zugänge freigibt. Je mehr Autonomie und Datenzugriff ein Modell erhält, desto wichtiger wird es zu verstehen, welche Quellen es nutzt, welche Schritte es ausführt und wo Fehler entstehen können.

Ein praktischer, oft unterschätzter Faktor ist schließlich das begrenzte *Kontextfenster* eines Modells. Das beschreibt die Menge an Text, die es in einem Gespräch gleichzeitig „im Blick“ behalten kann. Anders als ein menschlicher Berater verfügt ein Sprachmodell nicht automatisch über ein verlässliches Gedächtnis über einzelne Sitzungen hinweg. Wer einen neuen Chat beginnt, sitzt gewissermaßen einem neuen Berater gegenüber, der die eigene Vorgeschichte nicht kennt. Für eine einmalige Frage spielt das kaum eine Rolle. Wer die KI aber fortlaufend nutzt, etwa als Haushalts- oder Finanzplaner, erhält ohne ausreichenden Kontext nur allgemeine statt wirklich auf die eigene Lage zugeschnittene Ratschläge. Alternativ muss der Nutzer bei jeder Sitzung alle relevanten Eckdaten erneut und vollständig schildern, wobei leicht genau die Information fehlt, auf die es ankäme. Im schlimmsten Fall trifft das Modell dann stillschweigend Annahmen, die plausibel klingen, aber zur tatsächlichen Lebenssituation nicht mehr passen.

Hinzu kommt, dass Modelle Zeit nicht wie ein menschlicher Berater verstehen. Wird ein Chat ein Jahr später fortgesetzt, prüft das Modell nicht automatisch, ob sich Einkommen, Familienstand, Risikoneigung oder Anlageziel verändert haben. Selbst innerhalb eines langen Gesprächs können frühe Angaben zudem aus dem Kontextfenster fallen und dadurch praktisch in Vergessenheit geraten. Wer ein Modell ernsthaft für die eigene Finanzplanung einsetzt, sollte die relevanten Informationen daher bewusst und vollständig aktualisieren, statt sich darauf zu verlassen, dass die KI sich an Früheres verlässlich „erinnert“.

3.2 Überzeugungskraft und Aversion

Die Eloquenz moderner Sprachmodelle ist kein Zufall. Durch Trainingsverfahren, die Modelle an menschlichen Bewertungen ausrichten, insbesondere bestärkendes Lernen aus menschlichem Feedback (englisch *Reinforcement Learning from Human Feedback*, kurz *RLHF*), werden sie darauf optimiert, Antworten zu erzeugen, die Menschen als hilfreich, plausibel und überzeugend empfinden. Anders als bei einem menschlichen Gegenüber fehlen dabei jedoch die nonverbalen Signale, an denen sonst Unsicherheit oder Täuschung auffallen können. Kein ausweichender Blick, kein Stocken, keine nervösen Wiederholungen. Stattdessen liefert das Modell durchgehend flüssigen Text ohne Füll- und Verlegenheitswörter. Gerade diese sprachliche Glätte macht seine Aussagen überzeugend, selbst wenn sie unbelegt, frei erfunden oder faktisch falsch sind; solche plausibel klingenden, aber falschen Ausgaben werden als *Halluzinationen* bezeichnet.

Im Finanzkontext kann das bedeuten, dass ein Modell eine plausibel klingende, aber frei erfundene Quartalszahl nennt, ein wörtliches Zitat aus einem Geschäftsbericht fabriziert oder eine nicht existierende Kennziffer angibt. Für den Nutzer verschwimmt dadurch die Grenze zwischen fundierter Analyse und bloßer rhetorischer Überzeugungskraft. Das begünstigt ein übermäßiges Vertrauen (*Overreliance*) und die *Kontrollillusion*: Weil der Anleger den Dialog durch eigene Prompts steuert, glaubt er, Herr des Verfahrens zu sein, und übersieht dabei leicht, wie sehr die maschinelle Aufbereitung seine Wahrnehmung lenken kann.

Finanzakteure bewegen sich zwischen zwei Polen: unreflektierter Unterwerfung unter die Maschine und irrationaler Totalablehnung.

Diesem Übervertrauen steht ein ebenso problematischer Gegenpol gegenüber. Sobald die scheinbar souveräne Maschine einen sichtbaren Fehler macht, kann das Vertrauen abrupt kippen. Hier greift die *Algorithmenaversion*, die Dietvorst et al. (2015) untersuchen. Ihr kontraintuitiver Befund: Menschen verlieren nach einem beobachteten Fehler schneller das Vertrauen in einen Algorithmus als in einen Menschen und meiden ihn anschließend selbst dann, wenn er im Durchschnitt nachweislich treffsicherer ist. Das Bild ist allerdings differenzierter, als es zunächst scheint. Lässt man Menschen den Algorithmus auch nur geringfügig anpassen, steigt ihre Akzeptanz wieder deutlich (Dietvorst et al., 2018). In anderen Kontexten zeigt sich sogar das Gegenteil, eine *Algorithmus-Wertschätzung*: Laien folgen einem Rat eher, wenn er von einem Algorithmus zu stammen scheint (Logg et al., 2019).

Aversion und Wertschätzung sind also keine Widersprüche, sondern zwei mögliche Reaktionen auf dieselbe Grundsituation. Entscheidend ist, ob Nutzer Fehler, Kontrolle und eigene Verantwortung richtig einordnen. Gerade im Finanzkontext ist deshalb weder blindes Vertrauen noch pauschale Ablehnung angemessen, sondern ein reflektierter Umgang mit maschinellen Empfehlungen.

3.3 Korrektiv und Debiasing

Versteht man diese Mechanismen, lässt sich generative KI auch als Werkzeug der kognitiven Entzerrung (*Debiasing*) einsetzen. Gerade weil ein Sprachmodell keine eigene Karriere und keine emotionale Bindung an eine Anlageidee hat, kann es helfen, blinde Flecken sichtbar zu machen. Ein wirksamer Ansatz ist es, die KI als *Advocatus Diaboli*, also als bewussten Gegenargumentierer, einzusetzen. Ein Portfoliomanager speist seine optimistische Investmentthese ein und weist das Modell an: „Agiere als skeptischer Leerverkäufer und zerlege diese These Punkt für Punkt.“ Die KI liefert in Sekunden eine strukturierte Gegenargumentation, und das ohne die zwischenmenschlichen Hürden eines echten Widerspruchs: Sie wird nicht müde, nimmt nichts persönlich, kennt keine Hierarchie und widerspricht, ohne dass sich jemand unbeliebt machen muss.

Ein verwandtes Verfahren ist die *Pre-Mortem-Analyse* gegen Selbstüberschätzung. Vor einer großen Entscheidung lässt man die KI ein Szenario entwerfen, in dem die Strategie einige Zeit später krachend gescheitert ist, und die wahrscheinlichsten Gründe dafür benennen.

Aus dem Risiko wird ein Werkzeug, wenn der Mensch die Maschine nicht als Orakel behandelt, sondern als kritischen Sparringspartner.

Entscheidend ist damit nicht nur, *was* man die KI fragt, sondern welche Rolle man ihr zuweist. Wer sie um Bestätigung bittet, erhält leicht Bestätigung. Wer sie systematisch auf Gegenargumente, Schwachstellen und Alternativszenarien ansetzt, kann genau jene Verzerrungen abschwächen, die Finanzentscheidungen sonst anfällig machen. Generative KI ersetzt damit kein eigenes Urteil, kann aber helfen, dieses Urteil besser zu prüfen.



04

Systemische und makro- ökonomische Implikationen

Die Folgen generativer KI reichen über individuelle Nutzung und psychologische Verzerrungen hinaus. Sobald Sprachmodelle in Finanzinstituten Informationen auswerten, Entscheidungen vorbereiten und operative Prozesse prägen, werden sie Teil der Infrastruktur, auf der Märkte und Geschäftsmodelle aufbauen. Damit entstehen nicht nur Effizienzgewinne, sondern auch neue Abhängigkeiten und Risiken. Dieses Kapitel betrachtet zunächst die Effizienz der Märkte, dann den Struktur- und Jobwandel der Branche, danach kriminellen Missbrauch und schließlich das Risiko, wenn viele Akteure dieselben wenigen Modelle nutzen.

4.1 Markteffizienz und Trends im Asset Management

Wenn Sprachmodelle unstrukturierte Texte in Echtzeit auswerten, schließt sich das Zeitfenster für Informationsvorsprünge. Geschäftsberichte, Analystencalls, Pressemitteilungen, Nachrichten und regulatorische Dokumente lassen sich heute in Sekunden zusammenfassen, vergleichen und in handelbare Hypothesen übersetzen. Die Folge ist ein beschleunigter Alpha Decay: Die Möglichkeit, allein durch schnelleres Lesen oder einfachere Textauswertung eine Überrendite zu erzielen, schwindet.

Dass solche Vorsprünge verfallen, ist belegt, aber kein Automatismus. Eine Studie von Jacobs und Müller (2020), an der einer von uns beteiligt ist, zeigt für Hunderte von Anomalien über 39 Märkte ein differenziertes Bild: Verlässlich abgebaut werden die Überrenditen nach der Veröffentlichung vor allem im US-Markt, während sie sich international deutlich weniger zurückbilden. Die Autoren deuten dies als Hinweis auf segmentierte Märkte, in denen Grenzen der Arbitrage eine schnelle Korrektur verhindern. Auf KI- und nachrichtenbasierte Signale übertragen heißt das: Wo viele Akteure dieselben Quellen gleichzeitig auslesen, dürfte ein Vorsprung besonders schnell schmelzen. Wo Märkte kleiner, weniger liquide, sprachlich fragmentierter oder schwerer zugänglich sind, könnten Informationsvorteile dagegen länger bestehen bleiben.

Für das Asset Management verschiebt sich dadurch der Wettbewerb. Informationsbeschaffung wird zur Standard-Ware. Wer nur schneller Geschäftsberichte zusammenfasst oder Nachrichten präziser auswertet, wird daraus auf Dauer kaum einen nachhaltigen Vorteil ziehen, sobald dieselben Werkzeuge breit verfügbar sind. Wertvoller werden stattdessen proprietäre Daten, bessere Modellvalidierung, saubere historische Tests, robuste Risikokontrolle und die Fähigkeit, maschinelle Signale in ein überzeugendes Anlageurteil zu übersetzen. Profitables aktives Management bleibt möglich, muss seinen Vorteil aber stärker aus Daten, Infrastruktur und Urteilskraft ziehen.

Für die Geldanlage selbst zeichnet sich ein zwiespältiger Trend ab. Einerseits spricht vieles weiterhin für *passives Investieren*. Je mehr Marktteilnehmer ähnliche KI-Werkzeuge nutzen, desto schneller werden öffentlich verfügbare Informationen verarbeitet und desto schwerer wird es, den Markt dauerhaft zu schlagen. Andererseits können spezialisierte Häuser, die eigene Daten, Rechenleistung und Modellkompetenz geschickt kombinieren, zumindest zeitweise neue Informationsvorteile erzielen. KI-gestützte Strategien und ihre Rückrechnungen (*Backtests*) deuten solche Möglichkeiten an. Ein Versprechen für die Zukunft ist das ausdrücklich nicht, denn ein im Rückblick starkes Muster muss sich live nicht wiederholen.

Damit dürfte sich das Asset Management weiter polarisieren. Auf der einen Seite stehen sehr günstige, breit gestreute Indexprodukte, die für die meisten Anleger eine robuste Basis bleiben. Auf der anderen Seite stehen hochspezialisierte, datengetriebene Strategien, die versuchen, mit proprietären Informationen, technischer Infrastruktur und menschlichem Urteil dort Vorteile zu erzielen, wo öffentliche KI-Werkzeuge bereits zum Standard geworden sind. Besonders schwierig wird es für klassische aktive Produkte, die hohe Gebühren verlangen, aber im Kern nur allgemein verfügbare Informationen verarbeiten. Diese werden in Zukunft noch stärker unter Rechtfertigungsdruck geraten.

Diese Entwicklung rückt eine grundlegende Frage wieder in den Vordergrund. Wenn immer mehr Kapital passiv investiert wird und viele aktive Akteure ähnliche KI-Systeme nutzen, wer sorgt dann noch für effiziente Preise? Wahrscheinlich wird diese Aufgabe stärker bei einer kleineren Gruppe gut ausgestatteter aktiver Häuser liegen. Sie heben sich nicht mehr durch das bloße schnellere Lesen von Nachrichten ab, sondern durch Datenzugang, Rechenleistung, Modellkompetenz und Urteilkraft in mehrdeutigen Situationen. Markteffizienz muss darunter nicht zwangsläufig leiden. Sie dürfte aber stärker davon abhängen, dass weiterhin genügend Akteure mit unterschiedlichen Daten, Modellen und Einschätzungen aktiv nach Fehlbewertungen suchen.

4.2 Struktur- und Jobwandel

Mit der Automatisierung von Informationsarbeit verändern sich auch die benötigten Fähigkeiten im Finanzsektor. Der Wandel betrifft nicht nur einzelne Hilfstätigkeiten, sondern nach und nach ganze Berufsbilder. Dass Finanzberufe besonders exponiert sind, ist gut belegt. Eloundou et al. (2024) ordnen Finanztätigkeiten zu den am stärksten von Sprachmodellen betroffenen Berufen. Entscheidend ist dabei zunächst nicht, ob ein Beruf vollständig ersetzt wird, sondern wie viele seiner einzelnen Aufgaben sich durch Sprachmodelle unterstützen, beschleunigen oder automatisieren lassen.

Was das praktisch bedeutet, ist zwiespältig. Einerseits wirkt KI als *Verstärker*. In einer vielzitierten Feldstudie steigert ein KI-Assistent die Produktivität im Kundenservice um durchschnittlich 14 Prozent, bei den unerfahrensten Beschäftigten sogar um 34 Prozent (Brynjolfsson, Li & Raymond, 2025). Gerade weniger erfahrene Mitarbeiter profitieren offenbar davon, dass das System bewährte Formulierungen, Prozesswissen und implizite Erfahrung verfügbar macht. Andererseits ist mehr Output nicht automatisch besser. Setzen Aktienanalysten ein KI-Werkzeug ein, liefern sie zwar umfangreichere Berichte mit mehr Quellen, doch ihre Gewinnprognosen werden deutlich ungenauer, mit rund 59 Prozent mehr Fehlern; die zusätzliche Informationsfülle erschwert offenbar das klare Urteil (Xue et al., 2025). Das ist die verhaltensökonomische Pointe: Mehr Information ersetzt nicht das Urteil.

KI macht aus reiner Informationsbeschaffung eine Standardware. Knapp und wertvoll bleibt das menschliche Urteil in mehrdeutigen Lagen.

Besonders deutlich trifft der Wandel Einstiegsrollen. Eine Analyse von Lohnabrechnungsdaten durch Brynjolfsson, Chandar und Chen (2025) findet seit der breiten Verfügbarkeit generativer KI einen relativen Beschäftigungsrückgang von rund 16 Prozent bei den 22- bis 25-Jährigen in stark KI-exponierten Berufen, während die Beschäftigung erfahrener Kräfte stabil bleibt. Finanznahe Tätigkeiten wie Buchhaltung und Wirtschaftsprüfung zählen dabei zu den am stärksten exponierten Bereichen.

Daraus erwächst ein struktureller Zielkonflikt. KI übernimmt zuerst die einfachen, gut beschreibbaren Aufgaben, also genau jene Einstiegstätigkeiten, an denen Berufsanfänger bislang ihr Handwerk lernen. Wenn Banken, Beratungen und Wirtschaftsprüfungsgesellschaften diese Stufen wegautomatisieren und entsprechend weniger Juniorinnen und Junioren einstellen, trocknet langfristig die Pipeline aus, aus der die erfahrenen Fachkräfte von morgen hervorgehen. Kurzfristig spart das Kosten, mittelfristig stellt sich aber die Frage, wie und wo künftige Senior-Expertise überhaupt noch entstehen soll.

Welche Arten von Tätigkeiten besonders betroffen sind, lässt sich ebenso genauer beobachten. Der *Anthropic Economic Index* schätzt nicht nur theoretisch, welche Berufe von KI betroffen sein könnten, sondern misst, wofür Menschen Sprachmodelle tatsächlich einsetzen (Anthropic, 2026). Vorne liegen Software-Entwicklung, Kundenservice, Datenerfassung sowie schreib- und analyselastige Büroberufe, zu denen auch viele Finanztätigkeiten gehören. Die Quintessenz ist differenzierter, als Schlagzeilen vermuten lassen. Hohe Exposition heißt bislang vor allem, dass KI viele Aufgaben eines Berufs verändert, nicht zwangsläufig, dass die Stelle vollständig verschwindet.

Zugleich mahnt die Makroökonomie zur Gelassenheit. Acemoglu (2025) schätzt den gesamtwirtschaftlichen Produktivitätsgewinn durch KI auf höchstens etwa 0,7 Prozent über zehn Jahre zusammengekommen. Sein Hauptargument ist, dass nur ein kleiner Teil aller Arbeitsaufgaben überhaupt für KI zugänglich ist und sich davon wiederum nur ein Bruchteil in absehbarer Zeit wirtschaftlich automatisieren lässt. KI ist damit ein bedeutender, aber kein über Nacht alles umstürzender Faktor.

Für den Finanzsektor heißt das: Routinemäßige Analysearbeit wird zunehmend zur ersetzbaren Standardware, während andere Fähigkeiten an Bedeutung gewinnen. Aufgewertet werden Urteilkraft in unklaren Situationen, der Umgang mit Kunden in Krisen („Behavioral Coaching“) und die Fähigkeit, maschinelle Ergebnisse kritisch zu prüfen. Der Finanzakteur der Zukunft ist damit weniger reiner Informationsverarbeiter als verhaltensökonomischer Navigator und „algorithmischer Auditor“.

Noch grundsätzlicher gewinnt an Wert, was sich gerade nicht automatisieren lässt: persönliche Netzwerke, Vertrauen und Reputation. Je mehr die reine Analyse zur Standardware wird, desto stärker verschiebt sich der Wettbewerb auf das, was ein Modell nicht liefern kann. Belastbare Beziehungen öffnen den Zugang zu Kapital, zu Mandaten und zu Informationen, die noch nicht in den Daten stehen. Sie tragen auch durch Krisen, in denen Verlässlichkeit mehr zählt als Rechenleistung. Wer im Finanzsektor künftig heraussticht, dürfte daher weniger derjenige sein, der Informationen am schnellsten beschafft, als derjenige, der Menschen zusammenbringt und über Jahre gewachsenes Vertrauen genießen darf.

Mit zunehmender Autonomie verschiebt sich zudem eine Grundsatzfrage: Wie stark muss der Mensch eingebunden bleiben (*Human in the Loop*)? Verantwortung lässt sich nicht einfach an ein Modell delegieren. Gleichzeitig ist jede menschliche Freigabe teuer und langsam. Funktioniert die KI eigenständig gut, kann ein Institut, das stets einen Menschen zwischenschalten muss, gegenüber automatisierten Wettbewerbern ins Hintertreffen geraten. Die eigentliche Frage ist daher weniger, ob ein Mensch beteiligt ist, sondern wie viel Autonomie man der Maschine in welcher Situation einräumt.

Was die Aufsicht dazu sagt, ist eindeutiger als oft angenommen. Für Hochrisiko-Anwendungen verlangt der EU AI Act eine wirksame menschliche Aufsicht, und Aufseher wie BaFin und EZB betonen, dass die Verantwortung beim beaufsichtigten Institut bleibt, ganz gleich, wie viel die Maschine vorbereitet hat. Ein Mensch muss also nicht jede Entscheidung selbst treffen, aber er muss sie verstehen, überstimmen und verantworten können.

Hinzu kommen technische und geopolitische Abhängigkeiten, die den Strukturwandel begrenzen können. Wer zentrale Arbeitsprozesse auf wenige externe Modelle, Cloud-Anbieter oder Programmierschnittstellen stützt, macht sich von deren Verfügbarkeit abhängig. Fällt eine Schnittstelle aus, ändert ein Anbieter seine Nutzungsbedingungen oder wird der Zugang aus politischen Gründen beschränkt, betrifft das nicht mehr nur die IT-Abteilung, sondern unmittelbar die Arbeitsfähigkeit ganzer Teams. Auch Chip-Exportkontrollen, nationale Zugangsregeln oder die Abschaltung einzelner Modelle können damit zu einem operativen Risiko für Finanzinstitute werden. Je stärker KI in den Arbeitsalltag eingebettet wird, desto mehr wird technologische Resilienz selbst zu einer Managementaufgabe.

Ein eigenes, übergreifendes Risiko bleibt schließlich der Black-Box-Charakter vieler Modelle. Warum ein System zu einer bestimmten Einschätzung kommt, lässt sich oft nur schwer nachvollziehen, was Aufsicht, Haftung und interne Qualitätskontrolle erschwert. Zugleich entstehen neue Berufsbilder, die genau hier ansetzen, etwa in der KI-Regulierung, Cybersicherheit, Daten- und Qualitätskontrolle sowie in der Governance solcher Systeme. Die Finanzbranche dürfte damit einer der ersten „Kanarienvögel“ dafür sein, wie generative KI den Arbeitsmarkt verändert. Das liegt vor allem daran, dass der Finanzsektor seit jeher datengetrieben arbeitet und neue Technologien schnell in Prozesse, Produkte und Geschäftsmodelle übersetzt.

4.3 Marktmanipulation und Betrugsszenarien

Generative KI senkt nicht nur die Hürden für normale Privatanleger, sondern auch für Kriminelle. Sie skaliert Betrug in bisher unerreichter Geschwindigkeit und Qualität. Das eindrucklichste Beispiel sind audiovisuelle *Deepfakes*: 2024 verleiteten Betrüger in Hongkong einen Finanzmitarbeiter des Ingenieurkonzerns Arup in einer gefälschten Videokonferenz dazu, rund 25,6 Millionen US-Dollar zu überweisen (Fortune, 2024). Bemerkenswert ist, dass hier kein offensichtlich leichtgläubiges Opfer hereinfiel. Der Mitarbeiter schöpfte zunächst Verdacht und hielt die Aufforderung für einen Betrugsversuch, ließ sich dann aber vom täuschend echten Videocall mit vermeintlichen Kolleginnen und Vorgesetzten überzeugen. Das ist die unbequeme Lehre: Bei heutiger Fälschungsqualität kann es praktisch jeden treffen, der nicht über einen zweiten Kanal gegenprüft.

Im deutschsprachigen Raum kursieren massenhaft Deepfake-Werbevideos und gefälschte Anzeigen mit den geklonten Identitäten Prominenter. Allein für den Investor Frank Thelen wurden binnen weniger Monate rund 1.000 gefälschte Werbeanzeigen auf Meta-Plattformen gezählt (Business Insider, 2025); gefälschte „Tagesschau“-Beiträge und KI-Videos missbrauchten unter anderem Carsten Maschmeyer (dpa-Faktencheck, 2025), Christian Lindner mit der Plattform „Green-Vest“ (Stiftung Warentest, 2024) und Friedrich Merz (Handelsblatt, 2024). Wie real der Schaden ist, zeigt ein Fall aus Heilbronn. Ein Ehepaar überwies über mehrere Monate hinweg mehr als 400.000 Euro an eine Krypto-Plattform, die mit dem missbräuchlich verwendeten Namen und Bild von Günther Jauch beworben wurde (Stuttgarter Zeitung, 2025).

Solche Maschen zielen längst nicht nur auf Anlagebetrug. Besonders perfide sind *Schockanrufe* mit geklonten Stimmen: Aus wenigen Sekunden Audio, oft aus sozialen Netzwerken, erzeugen Täter die täuschend echte Stimme eines Angehörigen. Am Telefon meldet sich dann scheinbar die eigene Tochter unter Tränen, ein angeblicher Polizist übernimmt und verlangt nach einem angeblich schweren Unfall sofort eine Kautions. Verbraucherschützer warnen seit 2025/26 ausdrücklich vor dieser KI-Variante des „Enkeltricks“ (Verbraucherzentrale, 2026). Bei einzelnen Betroffenen führte sie bereits zu Schäden im fünf- bis sechststelligen Bereich. Weil die Masche so vertraut klingt, trifft sie Menschen gerade dort, wo sie es am wenigsten erwarten.

Das Ausmaß ist erheblich. Das FBI verzeichnete für 2025 Rekordverluste durch Internetkriminalität. Die gemeldeten Schäden lagen bei mehr als 20 Milliarden US-Dollar, mit Anlagebetrug als größter Einzelkategorie (Federal Bureau of Investigation, Internet Crime Complaint Center, 2025). Europol beschreibt KI im *EU Serious and Organised Crime Threat Assessment* als Beschleuniger organisierter Kriminalität (Europol, 2025). Die Werkzeuge reichen vom geklonten Stimmenanruf über synthetische Identitäten bis zu KI-generierten Phishing-Seiten und gefälschten „KI-Trading-Systemen“. In der „Operation Herakles“ nahmen deutsche und europäische Behörden 2025 in zwei Wellen über 2.200 betrügerische Handelsplattform-Domains vom Netz, die teils mit KI „am Fließband“ produziert worden waren (Bundesanstalt für Finanzdienstleistungsaufsicht, 2025).

Aus verhaltensökonomischer Sicht ist Betrug keine rein technische, sondern vor allem eine psychologische Attacke. Er zielt oft auf das *Autoritätsprinzip* (eine bekannte Stimme oder ein prominentes Gesicht setzt kritisches Prüfen aus), auf künstliche Dringlichkeit (wenn angeblich sofort gehandelt werden muss) und auf den *Herdentrieb* (etwa durch gefälschte Communities in Messenger-Gruppen). Wie anfällig selbst aktive Anleger sind, zeigt eine Befragung der FINRA-Stiftung aus dem Jahr 2025 (FINRA Investor Education Foundation, 2025): Die Hälfte der Befragten würde in ein hypothetisches Angebot mit „garantierter, risikofreier“ Rendite von 25 Prozent pro Jahr investieren. Eine solche sichere Rendite gibt es, ausgenommen in Zeiten von Hyperinflation, seriös nicht.

So prüfen Anleger einen Anbieter

Vor jeder Einzahlung lässt sich kostenlos prüfen, ob ein Anbieter überhaupt zugelassen ist:

- **Deutschland:** BaFin-Unternehmensdatenbank für beaufsichtigte Institute und BaFin-Warnmeldungen; ergänzend Handels- und Unternehmensregister für die Existenz der Firma.
- **Österreich:** FMA-Unternehmensdatenbank und FMA-Investorenwarnungen.
- **Schweiz:** FINMA-Warnliste sowie die Liste der bewilligten Institute.
- **EU/international:** ESMA-Register und IOSCO-Warnportal.

Vorsicht auch bei *regulatorischen Klonen*: Betrüger geben sich als eine echte, lizenzierte Firma aus, nutzen deren Namen und Registernummer, aber haben eine eigene Website, Telefonnummer und Bankverbindung. Ein Registertreffer belegt dann nur den *gestohlenen* Namen. Verifizieren Sie Kontaktdaten daher stets über das offizielle Register, nicht über die vom Anbieter genannten Angaben.

Die wirksamste Verteidigung ist Skepsis: Garantierte Renditen, Zeitdruck und Geheimhaltung sind die drei verlässlichsten Warnsignale.

Zum Schutz helfen zudem folgende Regeln:

- **Zweiten Kanal nutzen.** Jede unerwartete Zahlungsaufforderung sollte über einen selbst gewählten Kanal verifiziert werden, etwa durch einen Rückruf unter einer bekannten Nummer. In Familien und Unternehmen kann zusätzlich ein vorher vereinbartes Kennwort helfen.
- **Zeitdruck widerstehen.** Betrüger erzeugen bewusst Dringlichkeit. Wer sofort überweisen, investieren oder schweigen soll, sollte gerade deshalb innehalten.
- **Echtheit nicht nach Gefühl beurteilen.** Deepfakes sind heute nicht mehr zuverlässig „mit bloßem Auge“ zu erkennen. Ein überzeugendes Video, eine vertraute Stimme oder ein professioneller Internetauftritt ersetzen keine unabhängige Prüfung.
- **Empfängernamen prüfen.** Seit Oktober 2025 müssen Banken und Zahlungsdienstleister im Euroraum vor der Freigabe abgleichen, ob der angegebene Empfängername zur IBAN passt (*Verification of Payee*) (Europäische Kommission, 2025). Die Warnung „Name stimmt nicht überein“ ist ein deutliches Stoppsignal, auch wenn sie keine vollständige Betrugsprüfung ersetzt.

Die Technik kann jedoch auch Teil der Verteidigung sein. Sprachmodelle helfen dabei, Hinweise zu ordnen, Warnlisten zu finden und Anbieter, Websites oder Zahlungsaufforderungen kritisch einzuordnen. Verlässlich ist das nicht, aber als zusätzliche Prüfstufe nützlich. Die wichtigste Konsequenz bleibt deshalb einfach: Was echt aussieht, muss nicht echt sein. Bei wichtigen oder eiligen Entscheidungen braucht es immer eine unabhängige Bestätigung, etwa über einen persönlichen Rückruf oder ein vorab vereinbartes Codewort.

Auch der Gesetzgeber reagiert: Regulatorische Antwort – EU AI Act

Mit der KI-Verordnung (EU) 2024/1689 hat die Europäische Union einen einheitlichen Rechtsrahmen für Künstliche Intelligenz geschaffen. Für die Finanzbranche besonders relevant: KI-Systeme zur Bewertung der Kreditwürdigkeit natürlicher Personen gelten grundsätzlich als *Hochrisiko-Systeme*; Systeme zur reinen Betrugserkennung sind davon ausgenommen. Für Hochrisiko-Systeme gelten unter anderem Anforderungen an Risikomanagement, Dokumentation, Transparenz und menschliche Aufsicht (Europäisches Parlament und Rat der Europäischen Union, 2024).

Artikel 50 verlangt zudem die technische Kennzeichnung bestimmter KI-generierter Inhalte und die Offenlegung von Deepfakes. Diese Pflichten gelten grundsätzlich ab dem 2. August 2026. Für Systeme, die bereits zuvor auf den Markt gebracht wurden, muss die technische Kennzeichnung nach Artikel 50 Absatz 2 spätestens ab dem 2. Dezember 2026 erfüllt werden. Die zentralen Anforderungen für eigenständige Hochrisiko-Systeme nach Anhang III, darunter Systeme zur Kreditwürdigkeitsprüfung, werden nach der 2026 beschlossenen Änderung ab dem 2. Dezember 2027 anwendbar (Europäisches Parlament und Rat der Europäischen Union, 2026).

4.4 Herding 2.0: Das Risiko einer Modell-Monokultur

Ein eigenes, potenziell unterschätztes Systemrisiko entsteht, wenn viele Marktteure dieselben wenigen Modelle nutzen. Schon die klassische Verhaltensökonomie kennt das *Herding*, also Herdenverhalten an den Märkten. Generative KI kann dieses Muster verschärfen. Wenn Tausende Anleger, Analysten und Institutionen ihre Einschätzungen aus denselben Sprachmodellen, denselben Datenquellen und ähnlichen Prompts ziehen, entstehen leichter korrelierte Urteile und gleichgerichtete Trades. Das Problem ist dann nicht mehr nur, dass einzelne Akteure Fehler machen. Das Problem ist, dass viele Akteure denselben Fehler zur selben Zeit machen könnten.

Dahinter steht ein zentraler Mechanismus funktionierender Märkte: Meinungsverschiedenheit. Märkte brauchen unterschiedliche Einschätzungen, weil nur dann Käufer und Verkäufer aufeinandertreffen. Für denselben Preis hält der eine Akteur eine Aktie für unterbewertet, der andere für überbewertet, genau daraus entstehen Handel, Liquidität und Preisfindung. Ziehen dagegen viele Akteure ähnliche Schlüsse aus denselben Modellen, verengt sich die Meinungsbildung. Positionen können einseitiger werden, Liquidität kann gerade in Stressphasen verschwinden, und Stimmungsumschwünge können abrupt ausfallen.

Dass Algorithmen ohne ausdrückliche Absprache zu gleichförmigem Verhalten finden, ist nicht neu. Calvano et al. (2020) zeigen, dass lernende Preisalgorithmen eigenständig kollusive, also stillschweigend abgestimmte Muster entwickeln können. Für Sprachmodelle deuten Marktsimulationen in eine ähnliche Richtung. LLM-Agenten können in simulierten Märkten konsistente Handelsrollen einnehmen, Preisfindung, Blasen und strategische Liquiditätsbereitstellung erzeugen und durch gleichartige Eingaben korrelierte Entscheidungen treffen (Lopez-Lira, 2025). Für reale Märkte ist das noch kein Beweis für ein neues Krisenszenario, aber ein Hinweis darauf, dass Modell- und Promptvielfalt für Marktstabilität relevant werden können.

Die algorithmische Vergangenheit zeigt, wie schnell technische Gleichläufigkeit die Marktmechanik belasten kann. Beim „Flash Crash“ vom Mai 2010 brachen US-Indizes binnen Minuten massiv ein und erholten sich kurz darauf wieder. Damals standen klassische Handelsalgorithmen, automatische Ausführungsregeln und austrocknende Liquidität im Mittelpunkt. Künftig könnte sich ein Teil dieses Gleichschritts von der reinen Ausführung auf die Ebene der Analyse verlagern. Wenn viele Akteure nicht nur ähnlich handeln, sondern bereits zuvor ähnliche maschinelle Einschätzungen erhalten, wird Vielfalt der Modelle, Datenquellen und menschlichen Urteile zu einer Frage der Finanzstabilität.

Hinzu kommt das Risiko gemeinsamer Infrastruktur. Wenn sich Finanzinstitute auf wenige Modellanbieter, Cloud-Plattformen oder Programmierschnittstellen stützen, entsteht ein neues Klumpenrisiko. Dieses Risiko kann prozyklisch wirken. In ruhigen Zeiten funktionieren die Systeme scheinbar reibungslos, in Stressphasen hingegen greifen besonders viele Marktteilnehmer gleichzeitig auf dieselben Modelle zu und benötigen schnelle Antworten. Genau dann können Überlastung, Ratenlimits, verzögerte Antworten oder technische Ausfälle besonders folgenreich werden. Aus einem Effizienzgewinn im Normalbetrieb wird so ein ernstzunehmendes operatives Risiko im Krisenfall.

Diese Abhängigkeit ist nicht nur technischer, sondern auch geopolitischer Natur. Der Zugang zu leistungsfähigen Modellen, Chips, Cloud-Infrastruktur und Schnittstellen kann durch Exportkontrollen, nationale Regulierung, Anbieterentscheidungen oder politische Konflikte eingeschränkt werden. Wie schnell daraus ein handfestes Zugangsrisiko werden kann, zeigte sich im Juni 2026: Die US-Regierung ordnete an, den Zugang ausländischer Nutzer zu Anthropic-Modellen wie Fable 5 und Mythos 5 zu beschränken; weil sich Staatsangehörigkeit in Echtzeit kaum zuverlässig trennen ließ, schaltete Anthropic die betroffenen Modelle vorübergehend für alle Nutzer ab (CNN, 2026). Für Finanzinstitute, Forschungseinrichtungen und andere Anwender außerhalb der USA wird damit sichtbar, dass ein Modellzugang nicht nur eine technische, sondern auch eine politische Voraussetzung ist. Was im Alltag wie ein gewöhnliches Softwarewerkzeug erscheint, kann in angespannten Marktphasen zur kritischen Infrastruktur werden.

Systemisch relevant ist also nicht nur, was ein einzelnes Modell empfiehlt, sondern wie viele Akteure davon abhängig sind. Eine falsche Einschätzung bleibt handhabbar, solange sie vereinzelt auftritt und andere Marktteilnehmer dagegenhalten. Gefährlicher wird sie, wenn sie durch ein dominantes Modell, eine weit verbreitete Schnittstelle oder ein gemeinsames Datenproblem gleichzeitig in viele Entscheidungen eingeht. Dann wird aus Modellrisiko ein Marktstrukturproblem.

Die naheliegende Antwort ist nicht, KI aus den Märkten herauszuhalten. Dafür ist der Effizienzgewinn zu groß und die Entwicklung zu weit fortgeschritten. Wichtiger ist Resilienz. Es braucht unterschiedliche Modelle, unterschiedliche Datenquellen, unabhängige Plausibilitätsprüfungen, menschliche Eskalationswege und technische Ausweichlösungen. Gerade weil generative KI Märkte schneller machen kann, muss die Branche darauf achten, dass sie diese nicht zugleich gleichförmiger und störungsanfälliger macht.



Zusammenfassung und Schlusswort

Die Integration generativer KI in den Finanzsektor ist mehr als der nächste Effizienzsprung. Der Graben zwischen quantitativer Datenanalyse und der Auswertung weicher, sprachlicher Informationen schließt sich. Geschäftsberichte, Nachrichten, Beratungsgespräche, Research-Dokumente und regulatorische Texte werden maschinell lesbar, vergleichbar und in Sekunden verwertbar. Diese Veränderung erfasst die gesamte Wertschöpfungskette, vom dialogbasierten Robo-Advisor über den virtuellen Assistenten in der Beratung bis zu neuen Methoden der Finanzforschung. Richtig eingesetzt, kann KI damit Produktivität, Transparenz und Zugänglichkeit im Finanzbereich merklich erhöhen.

Doch die Automatisierung entbindet den Menschen nicht von seinen verhaltensökonomischen Schwächen. Im Gegenteil: Gerade weil Sprachmodelle flüssig, plausibel und scheinbar souverän antworten, können sie Übervertrauen, Bestätigungsfehler und Kontrollillusionen verstärken. Eine überzeugend formulierte Antwort ist kein Beleg für Wahrheit. Ein Modell kann helfen, eine Entscheidung besser zu durchdenken, aber es kann sie auch in die falsche Richtung lenken, wenn Frage, Kontext oder Datenbasis verzerrt sind.

Auf der Systemebene treten weitere Risiken hinzu. KI beschleunigt den Verfall einfacher Informationsvorsprünge, rüstet Kriminelle mit skalierbaren Werkzeugen für Betrug und Manipulation aus und kann durch gleichförmige Modellnutzung neue korrelierte Risiken schaffen. Zugleich kann der Finanzsektor abhängiger von wenigen Modellen, Anbietern, Chips, Schnittstellen und Datenquellen werden. Was als gewöhnliches Softwarewerkzeug beginnt, wird zur kritischen Infrastruktur werden. Wer wichtige Prozesse allein auf ein einzelnes Spitzenmodell stützt, sollte deshalb technische, regulatorische und geopolitische Abhängigkeiten von Anfang an mitdenken.

Bei aller Automatisierung bleibt der zentrale Grundsatz bestehen, dass Verantwortung sich nicht an ein Modell delegieren lässt. Die entscheidende Frage lautet daher nicht, ob Mensch oder Maschine überlegen ist, sondern wie beide sinnvoll zusammenspielen. Wahrscheinlich wird sich daher ein hybrides Modell durchsetzen. Die Maschine liefert Entwürfe, Gegenargumente, Zusammenfassungen und Warnsignale, während der Mensch prüft, gewichtet, verantwortet und entscheidet. Besonders wertvoll bleiben dabei (Soft-)Skills, die sich nicht leicht automatisieren lassen, etwa Urteilskraft in mehrdeutigen Situationen, Vertrauen, Reputation, persönliche Beziehungen und die Fähigkeit, maschinelle Ergebnisse kritisch einzuordnen.

Nicht die KI ersetzt den Menschen im Finanzwesen. Aber der Mensch, der KI mit verhaltensökonomischem Verständnis kombiniert, ersetzt den, der die Entwicklung ignoriert.

Für die Praxis lassen sich daraus drei Leitplanken ableiten:

1. **KI als Sparringspartner nutzen, nicht als Orakel.** Gute Prompts suchen nicht nur Bestätigung, sondern Gegenargumente, Schwachstellen und Alternativszenarien. Wer die Maschine bewusst als kritischen Gegenpart einsetzt, kann eigene Verzerrungen reduzieren.
2. **Zentrale Aussagen unabhängig prüfen.** Eloquenz ist kein Beleg. Zahlen, Zitate, Quellen, Zulassungen und Zahlungsaufforderungen müssen bei wichtigen Entscheidungen über unabhängige Kanäle kontrolliert werden.
3. **Abhängigkeiten begrenzen.** Wer KI in wichtige Prozesse einbettet, sollte nicht nur das beste Modell wählen, sondern auch Ausweichmodelle, Datenkontrolle, lokale oder offene Alternativen und klare Notfallprozesse mitdenken.

Für Privatanleger gilt weiterhin, dass Rendite ohne Risiko ein Warnzeichen ist. Anbieter sollten vor jeder Einzahlung über offizielle Register und Warnlisten geprüft werden. Unerwartete Zahlungsaufforderungen gehören über einen zweiten, selbst gewählten Kanal verifiziert. Und bei KI-Antworten gilt: Je größer der finanzielle Schaden einer falschen Entscheidung wäre, desto wichtiger sind unabhängige Prüfung, zweite Bestätigung und gesunde Skepsis gegenüber allem, was zu sicher, zu dringend oder zu überzeugend klingt.

Mit Blick auf die nächsten Jahre ist weder die Euphorie angebracht, die in Werbeversprechen, sozialen Medien und manchen Branchenkommentaren rund um KI mitschwingt, noch die reflexhafte Ablehnung. Generative KI wird Routinearbeit übernehmen, Finanzwissen zugänglicher machen, Informationsvorsprünge verkürzen und neue Risiken schaffen. Unter dem Strich überwiegen für uns die Vorteile, sodass wir mit vorsichtigem Optimismus in die Zukunft blicken. Denn richtig eingesetzt, kann generative KI den Finanzbereich produktiver, transparenter und zugänglicher machen. Der große Gewinner ist der, der die Werkzeuge versteht, ihre Grenzen kennt und seine eigenen Verzerrungen ernst nimmt. Wir, die Autoren, hoffen, dass dieser Band ein Stück dazu beitragen konnte.

Transparenzhinweis zur KI-Nutzung

Ganz dem Motto dieses Bandes folgend haben wir generative KI zur Prüfung sprachlicher Formulierungen und zur Fehlerkorrektur genutzt. Die Inhalte haben wir jedoch selbst geprüft, ausgewählt und verantworten diese.

Literatur

Wissenschaftliche Quellen

- Acemoglu, D. (2025). The Simple Macroeconomics of AI. *Economic Policy*, 40(121), 13–58. <https://doi.org/10.1093/epolic/eiae042>
- Barber, B. M., Huang, X., Odean, T., & Schwarz, C. (2022). Attention-Induced Trading and Returns: Evidence from Robinhood Users. *The Journal of Finance*, 77(6), 3141–3190. <https://doi.org/10.1111/jofi.13183>
- Barber, B. M., & Odean, T. (2000). Trading Is Hazardous to Your Wealth: The Common Stock Investment Performance of Individual Investors. *The Journal of Finance*, 55(2), 773–806. <https://doi.org/10.1111/0022-1082.00226>
- Breitung, C., & Müller, S. (2025). Global Business Networks. *Journal of Financial Economics*, 166, Artikel 104007. <https://doi.org/10.1016/j.jfineco.2025.104007>
- Brynjolfsson, E., Chandar, B., & Chen, R. (2025). *Canaries in the Coal Mine? Six Facts about the Recent Employment Effects of Artificial Intelligence* (Working Paper). Stanford Digital Economy Lab. <https://digitaleconomy.stanford.edu/publications/canaries-in-the-coal-mine/>

- Brynjolfsson, E., Li, D., & Raymond, L. R. (2025). Generative AI at Work. *The Quarterly Journal of Economics*, 140(2), 889–942. <https://doi.org/10.1093/qje/qjae044>
- Calvano, E., Calzolari, G., Denicolò, V., & Pastorello, S. (2020). Artificial Intelligence, Algorithmic Pricing, and Collusion. *American Economic Review*, 110(10), 3267–3297. <https://doi.org/10.1257/aer.20190623>
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm Aversion: People Erroneously Avoid Algorithms After Seeing Them Err. *Journal of Experimental Psychology: General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2018). Overcoming Algorithm Aversion: People Will Use Imperfect Algorithms If They Can (Even Slightly) Modify Them. *Management Science*, 64(3), 1155–1170. <https://doi.org/10.1287/mnsc.2016.2643>
- Eloundou, T., Manning, S., Mishkin, P., & Rock, D. (2024). GPTs are GPTs: Labor Market Impact Potential of Large Language Models. *Science*, 384(6702), 1306–1308. <https://doi.org/10.1126/science.adj0998>
- Fama, E. F., & French, K. R. (2010). Luck versus Skill in the Cross-Section of Mutual Fund Returns. *The Journal of Finance*, 65(5), 1915–1947. <https://doi.org/10.1111/j.1540-6261.2010.01598.x>
- Gong, Z., & Müller, S. (2026). No matter your financial literacy: Simplicity wins when choosing a fund. *Journal of Behavioral and Experimental Finance*, 50, Artikel 101188. <https://doi.org/10.1016/j.jbef.2026.101188>
- Jacobs, H., & Müller, S. (2020). Anomalies across the Globe: Once Public, No Longer Existent? *Journal of Financial Economics*, 135(1), 213–230. <https://doi.org/10.1016/j.jfineco.2019.06.004>
- Kelly, B. T., Malamud, S., Schwab, J., & Xu, T. A. (2026). *Scaling Point-in-Time Language Models* (Working Paper Nr. w35247). National Bureau of Economic Research. <https://doi.org/10.3386/w35247>
- Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm Appreciation: People Prefer Algorithmic to Human Judgment. *Organizational Behavior and Human Decision Processes*, 151, 90–103. <https://doi.org/10.1016/j.obhdp.2018.12.005>
- Lopez-Lira, A. (2025). *Can Large Language Models Trade? Testing Financial Theories with LLM Agents in Market Simulations*. arXiv:2504.10789. <https://arxiv.org/abs/2504.10789>
- Lopez-Lira, A., & Tang, Y. (2026). Can ChatGPT Forecast Stock Price Movements? Return Predictability and Large Language Models [Accepted for publication]. *Journal of Financial Economics*. <https://doi.org/10.2139/ssrn.4412788>
- Loughran, T., & McDonald, B. (2011). When Is a Liability Not a Liability? Textual Analysis, Dictionaries, and 10-Ks. *The Journal of Finance*, 66(1), 35–65. <https://doi.org/10.1111/j.1540-6261.2010.01625.x>
- Tversky, A., & Kahneman, D. (1981). The Framing of Decisions and the Psychology of Choice. *Science*, 211(4481), 453–458. <https://doi.org/10.1126/science.7455683>
- Xue, J., Zhang, Q., & Zhu, W. (2025). *Generative AI for Analysts*. arXiv:2512.19705. <https://arxiv.org/abs/2512.19705>

Online-Quellen

- Anthropic. (2026). *The Anthropic Economic Index: Labor Market Impacts of AI*. <https://www.anthropic.com/research/labor-market-impacts>
- Bank for International Settlements. (2024). *Large Language Models: A Primer for Economists* (BIS Quarterly Review). Bank for International Settlements. https://www.bis.org/publ/qtrpdf/r_qt2412b.htm
- Bank of England & Financial Conduct Authority. (2024). *Artificial Intelligence in UK Financial Services – 2024*. Bank of England und FCA. <https://www.bankofengland.co.uk/report/2024/artificial-intelligence-in-uk-financial-services-2024>
- Bundesanstalt für Finanzdienstleistungsaufsicht. (2025). „Operation Herakles“: Schlag gegen Cyberkriminelle. https://www.bafin.de/SharedDocs/Veroeffentlichungen/DE/Pressemitteilung/2025/pm_2025_12_08_operation_herakles.html
- Business Insider. (2025). *Vorsicht, Fake! Wenn ein „Frank Thelen“ 50.000 Euro in vier Wochen verspricht*. <https://www.businessinsider.de/gruenderszene/media/vorsicht-fake-wenn-ein-frank-thelen-50-000-euro-in-vier-wochen-verspricht/>
- CNBC. (2025). *Goldman Sachs Is Piloting Its First Autonomous Coder in Major AI Milestone*. <https://www.cnbc.com/2025/07/11/goldman-sachs-autonomous-coder-pilot-marks-major-ai-milestone.html>
- CNN. (2026). *Anthropic Suspends Access to Its Most Powerful AI Models After U.S. Bars Foreign Nationals*. <https://www.cnn.com/2026/06/13/business/anthropic-mythos-model-national-security>
- dpa-Faktencheck. (2025). *Online-Betrug: Gefälschter Tagesschau-Artikel wirbt für dubiose Plattform*. <https://dpa-factchecking.com/germany/250815-99-742496/>
- eToro. (2025, Oktober). *Retail Investors Flock to AI Tools, with Usage Up 46% in One Year*. <https://www.etoro.com/news-and-analysis/etoro-updates/retail-investors-flock-to-ai-tools-with-usage-up-46-in-one-year/>
- Europäische Kommission. (2025). *New EU Rules Make Instant Euro Payments Faster and Safer*. https://finance.ec.europa.eu/news/new-eu-rules-make-instant-euro-payments-faster-and-safer-2025-10-10_en
- Europäisches Parlament und Rat der Europäischen Union. (2024). *Verordnung (EU) 2024/1689 zur Festlegung harmonisierter Vorschriften für künstliche Intelligenz (KI-Verordnung)* (Verordnung). Amtsblatt der Europäischen Union. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- Europäisches Parlament und Rat der Europäischen Union. (2026). *Regulation amending Regulations (EU) 2024/1689, (EU) 2018/1139 and (EU) 2023/1230 as regards the simplification of the implementation of harmonised rules on artificial intelligence (Digital Omnibus on AI)*. <https://data.consilium.europa.eu/doc/document/PE-30-2026-REV-1/en/pdf>
- Europol. (2025). *EU Serious and Organised Crime Threat Assessment (EU-SOCTA) 2025*. European Union Agency for Law Enforcement Cooperation. <https://www.europol.europa.eu/publication-events/main-reports/changing-dna-of-serious-and-organised-crime>

- Federal Bureau of Investigation, Internet Crime Complaint Center. (2026). *Internet Crime Report 2025*. U.S. Federal Bureau of Investigation. https://www.ic3.gov/AnnualReport/Reports/2025_IC3Report.pdf
- FINRA Investor Education Foundation. (2025). *Findings on Fraud Awareness Among Investors* (National Financial Capability Study). <https://www.finra.org/media-center/newsreleases/2025/finra-foundation-releases-findings-fraud-awareness-among-investors>
- Fortune. (2024, 17. Mai). *A Deepfake 'CFO' Tricked British Design Firm Arup in a \$25 Million Fraud*. <https://fortune.com/europe/2024/05/17/arup-deepfake-fraud-scam-victim-hong-kong-25-million-cfo/>
- Fox Business. (2025, 23. Juni). *Goldman Sachs Announces Firm-Wide Launch of AI Assistant*. <https://www.foxbusiness.com/technology/goldman-sachs-announces-firmwide-launch-ai-assistant>
- Google. (2025, 6. November). *New Google Finance AI Features: Deep Search and More*. <https://blog.google/products-and-platforms/products/search/new-google-finance-ai-deep-search/>
- Handelsblatt. (2024). *Deepfakes: Experten warnen vor gefälschten Werbevideos mit Promis*. <https://www.handelsblatt.com/technik/ki/deepfakes-experten-warnen-vor-gefaelschten-werbevideos-mit-promis/100064525.html>
- Harbert, T. (2021, Februar). *Tapping the Power of Unstructured Data* [Accessed: 2026-07-09]. <https://mitsloan.mit.edu/ideas-made-to-matter/tapping-power-unstructured-data>
- Morgan Stanley. (2023). *Key Milestone in Innovation Journey with OpenAI*. <https://www.morganstanley.com/press-releases/key-milestone-in-innovation-journey-with-openai>
- Reinsel, D., Gantz, J., & Rydning, J. (2018). *The Digitization of the World: From Edge to Core* (IDC White Paper). International Data Corporation. <https://www.seagate.com/files/www-content/our-story/trends/files/idc-seagate-dataage-whitepaper.pdf>
- Robinhood. (2026). *Robinhood Is Now Open to Agents*. <https://robinhood.com/us/en/newsroom/robinhood-is-now-open-to-agents/>
- S&P Dow Jones Indices. (2025). *SPIVA U.S. Scorecard: Year-End 2024*. S&P Dow Jones Indices LLC. <https://www.spglobal.com/spdji/en/documents/spiva/spiva-us-year-end-2024.pdf>
- Stiftung Warentest. (2024). *Green-Vest: TV-Beitrag mit Finanzminister gefälscht*. <https://www.test.de/Green-Vest-TV-Beitrag-mit-Finanzminister-gefaelscht-6119456-0/>
- Stuttgarter Zeitung. (2025). *400.000 Euro erbeutet: Krypto-Betrüger nutzen Bild von Günther Jauch*. <https://www.stuttgarter-zeitung.de/lokales/baden-wuerttemberg/400000-euro-erbeutet-krypto-betrueger-nutzen-bild-von-guenther-jauch-ehepaar-verliert-ersparnisse-78946231.html>
- Verbraucherzentrale. (2026, 26. Januar). *KI and me: Schockanrufe in neuem Gewand*. <https://www.verbraucherzentrale-brandenburg.de/pressemitteilungen/ki-and-me-schockanrufe-in-neuem-gewand-116706>

Wir über uns

Der Behavioral Finance e.V.

Die Forschungsrichtung der Behavioral Finance greift Erkenntnisse aus der Psychologie auf, um das Anlegerverhalten und andere Phänomene an den Kapitalmärkten zu erklären. Der Behavioral Finance e.V. hat es sich neben der Forschung zur Aufgabe gemacht, die finanzielle Entscheidungskompetenz in Deutschland zu erhöhen. Da die Sprache der Wissenschaft oftmals nicht einfach zu verstehen ist, möchten wir Ihnen die Kernaussagen und die neuesten und wichtigsten Erkenntnisse aus der Wissenschaft näherbringen. Dabei sollen insbesondere die Implikationen und Chancen für die Praxis, die sich aus diesen Forschungsergebnissen ergeben, herausgearbeitet und betont werden.

Autoren dieses Bandes

Prof. Dr. Sebastian Müller, CFA



Prof. Dr. Sebastian Müller, CFA ist Center Director des Centers for Digital Transformation und Inhaber der Professur für Finanzen an der Technischen Universität München. Seine Promotion schloss er 2011 an der Universität Mannheim ab. In seiner Forschung untersucht er empirisch die Preisbildung von Wertpapieren, Investorenverhalten und Trends der digitalen Finanzwirtschaft. Hierfür verwendet er nicht selten Methodiken des maschinellen Lernens oder generative KI.

M.Sc. David Osterrieder



M.Sc. David Osterrieder ist wissenschaftlicher Mitarbeiter und Doktorand an der Professur für Finanzen der Technischen Universität München. Zuvor studierte er Betriebswirtschaftslehre (Master) und Wirtschaftswissenschaften (Bachelor) an der Goethe-Universität in Frankfurt. In seinem Promotionsvorhaben beschäftigt er sich unter anderem mit der Anwendung von KI in der empirischen Finanzmarktforschung.

Change Log

Änderung

10.08.2026

Erste veröffentlichte Fassung