

KI, Verantwortung und psychologische Faktoren bei Investments

Interview mit Martin Nimbach

Oktober 2025

I. Einführung – Kontext und persönlicher Hintergrund

Frage: *Lieber Martin, vielen Dank, dass Du Dir Zeit nimmst für ein Interview zum Themenkomplex “KI und Ethik”! Magst Du zum Einstieg kurz erläutern wer Du bist und was Dein Beruf ist? In welcher Verbindung stehst Du zu KI und/oder ethischen Fragen?*

Martin Nimbach:

Ich komme aus der Ingenieurwissenschaft und habe viele Jahre an der Schnittstelle von Technologie und Organisation gearbeitet – in Konzernen, Startups und interdisziplinären Forschungsteams. Heute liegt mein Schwerpunkt auf der Verbindung von KI, Architektur, Führung und neurokognitiven Modellen – also Modellen, die beschreiben, wie Denken, Fühlen und Entscheiden im Gehirn entstehen. Sie verbinden Psychologie (Wahrnehmen, Denken, Fühlen, Entscheiden) mit Neurowissenschaft (Aktivität von Nervenzellen, Netzwerken, Neurotransmittern usw.).

Frage: *Was bedeutet KI für Dich?*

Martin Nimbach:

KI ist für mich ein neues, sehr flexibles Werkzeug. Damit lassen sich komplexe Themen bearbeiten, auch ohne traditionell zu programmieren. Große Sprachmodelle (Large Language Models; LLM) spiegeln viel von dem, was im Internet steht, und geben schnellen Zugang zu Wissen. Ihre Ergebnisse sind aber nicht unfehlbar – daher immer mit Vorsicht nutzen und nachprüfen.

II. Beobachtungen & Systemfragen

Frage: *Wie erlebst Du die aktuelle Debatte über KI – in Deiner Branche, Deinem Fachgebiet, Deinem Umfeld?*

Martin Nimbach:

In vielen Organisationen herrscht ein Spannungsfeld zwischen Hype, Überforderung und Fehlinformationen. KI wird häufig implementiert – aber nur selten wirklich *integriert*.

Frage: *Was wird Deiner Meinung nach übersehen oder unterschätzt?*

Martin Nimbach:

Die strukturelle Qualität der Modelle: Biases – das sind Verzerrungen, die die Ergebnisse beeinflussen, die eine KI liefert - Interpretierbarkeit und Governance. Viele Systeme sind schlicht nicht robust genug, um produktiv eingesetzt zu werden.

Frage: *Gibt es ethische Fragen, die Dich besonders beschäftigen?*

Martin Nimbach:

Ja – insbesondere die Rechtssicherheit bei der Nutzung der Daten, die als Input für das Training der Modelle dienen.

III. Verantwortung & Design

Frage: *Wer trägt Verantwortung für den ethischen Einsatz von KI – und wer übernimmt sie tatsächlich?*

Martin Nimbach:

Technisch gesehen braucht es klare Rollen: Audit-Instanzen, Ethik-Boards und eine belastbare Security-Architektur. Alles andere bleibt Wunschdenken.

Frage: *Was brauchen wir, um ethische Prinzipien wirksam zu machen?*

Martin Nimbach:

Verbindliche Standards, transparente Tests und verpflichtende Aufklärungspflichten bei kritischen Anwendungsfeldern.



Frage: *Hast Du eigene Prinzipien?*

Martin Nimbach:

Ja, unser Framework heißt „CLEAR“: Clarity – Leadership – Ethics – Awareness – Resilience.

IV. Praxis & Erfahrung

Frage: *Hast Du ein konkretes Beispiel für gute oder schlechte KI-Praxis?*

Martin Nimbach:

In einem Tech-Projekt wurde eine KI im HR-Screening eingesetzt – doch niemand prüfte die Trainingsdaten. Das führte zu systematischen Verzerrungen (Biases), die erst spät auffielen.

Frage: *Was ist „gute Praxis“?*

Martin Nimbach:

Eine nachvollziehbare Architektur, regelmäßige Audits, Explainability-by-Design (das heisst, dass die KI verständlich erklären kann, *wie* und *warum* sie zu einem bestimmten Ergebnis kommt), redundante Kontrollmechanismen (also mehrere, sich gegenseitig absichernde Systeme oder Verfahren) und ein durchdachtes Prompting. Das klingt technisch, heißt aber im Kern: KI soll nachvollziehbar und doppelt abgesichert sein.

V. Ausblick & Einstellung

Frage: *Was gibt Dir Sicherheit im Umgang mit KI?*

Martin Nimbach:

Der Einbau von regelbasierten Kontrollinstanzen, vergleichbar mit dem präfrontalen Kortex - dem Teil des Gehirns, der für Kontrolle und Abwägung zuständig ist. Unsere Intuition ähnelt einem Large Language Modell, doch die Natur hat uns eine Impulskontrolle gegeben, die als Filter vor unseren Handlungen wirkt. Ein entsprechendes Kontrollmodell mit ethischen Regeln sollte auch in KI-Systemen implementiert sein, bevor diese Ergebnisse ausgeben.

Frage: *Was bereitet Dir Sorgen?*



Martin Nimbach:

Ein unkritischer, naiver Einsatz von LLMs.

Frage: *Was muss in den nächsten fünf Jahren passieren?*

Martin Nimbach:

Wir brauchen verbindliche Sicherheitsstandards, ethische Architekturprinzipien und interdisziplinäre Auditmechanismen.

5 vertiefte Learnings für Kapitalgeber

1. Vertrauen ist keine KI-Eigenschaft – sondern eine Führungsleistung.

Vertrauen in KI entsteht nicht durch technische Leistungsfähigkeit, sondern durch verantwortungsbewusste Führung. Kein Modell kann Vertrauen erzeugen, wenn das Team nicht fähig ist, Entscheidungen transparent und verantwortungsvoll zu treffen. Es steht und fällt also mit der Art, wie *Menschen* mit KI arbeiten.

Investoren sollten fragen:

- Wer trägt die finale Verantwortung für Modellentscheidungen?
- Gibt es dokumentierte Entscheidungsprotokolle bei kritischen AI-Aussagen?
- Wird im Team aktiv widersprochen – und wie wird mit Dissens umgegangen?

2. Ein gutes Modell ohne Audit ist ein unkalkulierbares Asset.

Ein funktionierendes Modell allein reicht nicht – es muss erklärbar, testbar und auditierbar sein. Ohne Bias-Scans, Verifikationsschichten und externe Prüfprozesse entsteht strukturelle Intransparenz.

Investoren sollten fragen:

- Wie wird Bias erkannt, dokumentiert und reduziert?
- Wer führt Audits durch – intern, extern, regelmäßig?
- Gibt es ein Incident Response Playbook für Modellversagen?

3. Skalierbarkeit beginnt bei Governance – nicht bei GPU-Power.

Viele KI-Startups investieren in Modelltraining, vernachlässigen aber die Organisation, die es tragen muss. Skalierbare KI braucht skalierbare Governance – mit klaren Rollen, Versionierung und ethischer Architektur.

Investoren sollten fragen:

- Welche Rollen existieren für AI-Governance und Ethik?
- Gibt es Protokolle zur Nachverfolgung von Modellversionen und Verantwortlichkeiten?
- Ist das Unternehmen regulatorisch aufgestellt (z. B. im Hinblick auf den AI Act)?

4. Expertise ist nur dann wertvoll, wenn sie sich selbst hinterfragt.

Viele technische Teams sind exzellent, aber kognitiv homogen. Interdisziplinarität ist ein Sicherheitsfaktor. Nur vielfältige Teams mit psychologischer Sicherheit erkennen Schwachstellen frühzeitig – und lernen daraus.

Investoren sollten fragen:

- Wie divers ist das Team in Fachkultur, Perspektiven und Erfahrung?
- Wann wurde zuletzt externe Kritik einbezogen – und wie wurde reagiert?
- Gibt es interne Formate für strukturierte Fehleranalysen?

5. KI ist kein inhärenter Vorteil – sondern ein Verstärker.

KI verbessert nichts von sich aus – sie verstärkt das, was im System bereits angelegt ist. Gute Organisationen werden effizienter, dysfunktionale gefährlicher.

Investoren sollten fragen:

- Welche Schwächen oder Konflikte könnten durch KI verschärft werden?



- Dient die Technologie dazu, Klarheit zu schaffen – oder nur Effizienz zu steigern?
- Was hat das Unternehmen unternommen, um Struktur, Führung und Kultur KI-fit zu machen?

Diese fünf Perspektiven bieten Investoren eine ganzheitliche Bewertungsgrundlage: technisch belastbar, psychologisch fundiert und strukturell durchdacht. Sie helfen, nicht nur das Produkt, sondern die Organisation als Ganzes zu verstehen – und echte Resilienz von bloßer Performance zu unterscheiden.

Vielen Dank für Deine Zeit, lieber Martin!