

DIE STOCKFISH-AGENDA

Was übermenschliche KI für Kontrolle, Wettbewerb und
industrielle Wertschöpfung bedeutet

Lothar K. Doerr

Institut für Produktionserhaltung e. V.

August 2026

EINFÜHRUNG

Vom Schachbrett in die Wertschöpfung

Elon Musks Vergleich mit der Schachengine Stockfish beschreibt keine weitere Stufe der Automatisierung. Er markiert den Moment, in dem menschliche Fachleistung nicht mehr der oberste Maßstab ist.

Im Interview mit The Economist wählt Elon Musk für die erwartete Entwicklung künstlicher Intelligenz ein Bild aus dem Schach: KI werde beim Programmieren irgendwann das 'Stockfish level' erreichen. Dann gebe es für Menschen, wie Musk formuliert, 'just no way to compete'. Der Satz ist zugespitzt, aber analytisch ergiebig. Stockfish ist nicht bloß schneller als ein Großmeister. Die Software spielt in einer Leistungsklasse, in der ein direkter Wettbewerb mit Menschen seinen Sinn verloren hat.

Der Vergleich darf dennoch nicht zu wörtlich genommen werden. Schach ist eine geschlossene Welt. Regeln, Informationen und Ziel sind bekannt. Wirtschaftliche Entscheidungen bewegen sich dagegen in offenen Systemen: Ziele widersprechen sich, Daten sind unvollständig, Nebenwirkungen treten zeitversetzt auf. Gerade deshalb lohnt es sich, Musks Metapher weiterzudenken. Wenn KI in einzelnen Wissensdomänen tatsächlich eine strukturelle Überlegenheit erreicht, verschieben sich drei Grundlagen unternehmerischer Wertschöpfung.

Die Stockfish-Agenda beginnt dort, wo KI nicht mehr nur assistiert, sondern den menschlichen Leistungsmaßstab verlässt.

Die Metapher hat sechs wirtschaftliche Konsequenzen. Erstens verliert der Mensch in einzelnen Domänen seine Stellung als höchster Leistungsmaßstab. Zweitens wird aus technischer Beschleunigung eine strategische Zeitfrage. Drittens kann der Mensch selbst zum langsamsten Teil der Wertschöpfung werden. Viertens wird Kontrolle schwieriger, weil Ergebnisse beurteilt werden müssen, die niemand mehr aus eigener Leistung reproduzieren kann. Fünftens verliert der bloße Zugang zu leistungsfähiger KI seinen Seltenheitswert. Sechstens wächst die Bedeutung der physischen Welt: Software kann sich in Sekunden vervielfältigen, Fabriken, Stromnetze und Produktionslinien nicht.

Die folgenden sechs Beiträge untersuchen diese Verschiebungen. Sie verstehen Musks Fünfjahreshorizont nicht als gesicherte Prognose, sondern als strategisches Szenario. Die entscheidende Frage lautet nicht, ob das Stockfish-Level im Jahr 2031

pünktlich erreicht wird. Sie lautet, welche organisatorischen, wirtschaftlichen und industriellen Vorbereitungen heute erforderlich wären, falls die Entwicklung schneller verläuft als die Institutionen, die sie beherrschen sollen.

Stockfish-Level: Wenn der Mensch nicht mehr mithalten kann

Elon Musk greift auf das Schach zurück, weil dort bereits sichtbar ist, was strukturelle maschinelle Überlegenheit bedeutet: Der Mensch spielt weiter, aber er ist nicht länger die höchste Instanz.

Der Satz fiel beinahe beiläufig. Im großen Interview mit The Economist erklärte Elon Musk, künstliche Intelligenz sei beim Programmieren bereits besser als mindestens neunzig Prozent der professionellen Softwareentwickler. Bald werde sie besser als neunundneunzig Prozent sein. Danach, so Musk, gebe es 'just no way to compete'. Die Systeme erreichten dann das 'Stockfish level'.

Das Bild ist sorgfältiger gewählt, als es zunächst wirkt. Stockfish ist keine allgemeine künstliche Intelligenz, kein bewusstes Wesen und kein elektronischer Schachphilosoph. Es ist eine hochspezialisierte, frei verfügbare Open-Source-Engine. Gerade diese Begrenzung macht den Vergleich so aufschlussreich. Innerhalb eines klar umrissenen Feldes hat die Maschine einen Abstand aufgebaut, der nicht mehr durch Training, Talent oder einen besonders guten Tag des Menschen geschlossen werden kann.

Das Match, das nie stattfand

Stockfish hat Magnus Carlsen nie in einem offiziellen Wettkampf um die Weltmeisterschaft besiegt. Musks Formulierung beschreibt kein historisches Match, sondern ein hypothetisches Kräfteverhältnis. Der berühmte reale Wendepunkt liegt fast drei Jahrzehnte zurück: 1997 besiegte IBMs Deep Blue den amtierenden Weltmeister Garry Kasparow in einem Wettkampf unter regulärer Bedenkzeit mit 3,5 zu 2,5 Punkten. Es war die erste Matchniederlage eines amtierenden Weltmeisters gegen ein Computersystem.

Seitdem hat sich die Distanz dramatisch vergrößert. Magnus Carlsen führte die FIDE-Weltrangliste im Juli 2026 mit einer klassischen Elo-Zahl von 2.823 Punkten an. Stockfish 18 stand in der CCRL-Liste vom 23. Juli 2026 bei 3.649 Punkten. Diese Zahlen stammen aus unterschiedlichen Wertungssystemen und dürfen nicht wie Messwerte eines gemeinsamen Labors voneinander abgezogen werden. Sie zeigen dennoch die Größenordnung: Die stärkste Maschine spielt nicht knapp oberhalb des besten Menschen, sondern in einer eigenen Vergleichsklasse.

Schon die Stockfish-Entwickler schätzten 2020, eine damalige Version könne dem menschlichen Weltmeister bei klassischer Bedenkzeit einen Zeitvorteil von eins zu tausend gewähren und dennoch ungefähr auf Augenhöhe spielen. Inzwischen läuft eine nochmals deutlich stärkere Engine auf gewöhnlichen Rechnern und in angepasster Form auch auf Mobilgeräten. Musks Pointe lautet daher nicht, Carlsen sei schwach. Sie lautet, dass selbst das Außergewöhnliche des Menschen gegen maschinelle Skalierung irgendwann gewöhnlich wirkt.

Wie Stockfish stark wurde

Stockfish verbindet eine sehr schnelle Suche durch mögliche Zugfolgen mit einer Bewertungsfunktion, die aussichtsreiche von schlechten Stellungen unterscheidet. Seit 2020 nutzt die Engine NNUE, ein effizient aktualisierbares neuronales Netz. Es bewertet Positionen anhand sehr großer Mengen von Trainingsdaten, während die klassische Suchlogik weiterhin Varianten berechnet und vergleicht. Stockfish 18 trainiert seine Netze mit einem automatisierten Verfahren auf mehr als hundert Milliarden von Leela Chess Zero bewerteten Positionen.

Das Ergebnis ist nicht menschliches Denken in Silizium. Die Engine empfindet keinen Druck, fürchtet keine Niederlage und versteht nicht, was eine Weltmeisterschaft kulturell bedeutet. Sie sucht, bewertet und entscheidet unter genau definierten Bedingungen. Ihre Überlegenheit entsteht aus der Verbindung von Rechengeschwindigkeit, Datentiefe, statistischer Bewertung und der unermüdlichen Prüfung von Varianten.

Was Musks Vergleich tatsächlich besagt

Übertragen auf das Programmieren wäre das Stockfish-Level nicht erreicht, wenn eine KI gelegentlich besseren Code schreibt als ein Mensch. Es wäre erreicht, wenn der direkte Leistungsvergleich seinen praktischen Sinn verliert: wenn das System komplexe Software schneller entwirft, Fehler systematischer findet, Abhängigkeiten vollständiger überblickt und Lösungen hervorbringt, die der Mensch zwar prüfen, aber nicht mehr aus eigener Kraft verbessern kann.

Dann verschiebt sich die Rolle des Experten. Der Programmierer wäre nicht verschwunden, ebenso wenig wie Carlsen nach dem Siegeszug der Schachengines verschwunden ist. Aber er wäre nicht mehr selbstverständlich der beste Problemlöser im Raum. Seine Aufgaben lägen stärker in der Wahl des Problems, der Formulierung von Zielen, der Bewertung von Risiken, der Gestaltung von Systemgrenzen und der Verantwortung für Folgen.

Stockfish hat den Großmeister nicht überflüssig gemacht. Es hat ihm nur die Stellung als höchste Instanz genommen.

Gerade darin liegt die wirtschaftliche Bedeutung der Metapher. Unternehmen müssen sich auf eine Welt vorbereiten, in der Menschen weiterhin entscheiden und haften, obwohl sie die fachlich stärkste Einzelleistung nicht mehr selbst erbringen. Die Frage lautet dann nicht mehr, ob die Maschine intelligenter ist. Sie lautet, wer Ziele setzt, Ergebnisse freigibt und für Entscheidungen einsteht, deren Entstehung kein Mensch vollständig nachvollziehen kann.

Quellen und Datenbasis

The Economist: [The full-length interview with Elon Musk](#)

Stockfish: [Stockfish 18 - Release und technischer Hintergrund](#)

CCRL: [40/15 Rating List vom 23. Juli 2026](#)

FIDE: [Magnus Carlsen - offizielles Ratingprofil](#)

IBM: [Deep Blue und der Wettkampf gegen Garry Kasparow](#)

BEITRAG 07

Der Mensch bleibt verantwortlich. Aber versteht er noch, wofür?

Neuere Studien aus Schach, Medizin und Wissensarbeit zeigen, worin das eigentliche Stockfish-Level der Arbeitswelt bestehen könnte: nicht allein in der Überlegenheit der Maschine, sondern im Verlust der menschlichen Fähigkeit, diese Überlegenheit unabhängig zu prüfen.

Ein Schachgroßmeister kann einen Zug von Stockfish ablehnen. Die Regeln hindern ihn nicht daran. Er kann die Figur an eine andere Stelle setzen, eine andere Variante wählen, seinem Instinkt folgen. Nur wird er damit aller Wahrscheinlichkeit nach die Qualität der Stellung verschlechtern. Die formale Entscheidungsgewalt bleibt beim Menschen; die sachliche Autorität ist längst an die Maschine übergegangen.

Dieser Unterschied ist für die Arbeitswelt entscheidend. In den meisten Diskussionen über künstliche Intelligenz wird gefragt, ob die Maschine den Menschen ersetzt. Die neuere Forschung legt eine unbequemere Frage nahe: Was geschieht, wenn der Mensch zwar nicht ersetzt wird, aber nicht mehr in der Lage ist, die maschinelle Entscheidung aus eigener Kompetenz zu beurteilen?

Dann bleibt der Mensch im Organigramm verantwortlich, ohne noch die stärkste Instanz der Erkenntnis zu sein.

Die Maschine rechnet nicht nur mehr - sie hält mehr aus

Eine 2025 veröffentlichte und 2026 überarbeitete Untersuchung über Stockfish und Leela kommt zu einem bemerkenswerten Ergebnis.¹ Die Forscher verglichen Partien menschlicher Spitzenspieler mit Begegnungen führender Schachprogramme. Dazu entwickelten sie ein Maß für die strategische Spannung einer Stellung: Wie viele Angriffs-, Verteidigungs- und Kontrollbeziehungen bestehen gleichzeitig? Wie dicht ist das Netz der Abhängigkeiten? Wie weit können die Folgen eines einzelnen Fehlers durch die gesamte Stellung laufen?

Die stärksten Schachprogramme halten hochkomplexe, spannungsreiche Positionen länger aufrecht als selbst die besten menschlichen Spieler. Mit wachsender Rechentiefe nimmt diese Fähigkeit weiter zu. Die Programme

vereinfachen nicht vorschnell. Sie können mehr Unsicherheit, mehr gegenseitige Abhängigkeiten und mehr mögliche Fortsetzungen gleichzeitig beherrschen.

Das verändert die Bedeutung maschineller Überlegenheit. Stockfish ist nicht lediglich ein Großmeister mit höherer Geschwindigkeit. Das Programm operiert in einem Komplexitätsraum, den ein Mensch nur noch ausschnittsweise überblickt. Menschen vereinfachen schwierige Situationen nicht allein deshalb, weil Vereinfachung objektiv die beste Strategie wäre. Sie vereinfachen auch, weil Aufmerksamkeit, Gedächtnis, Rechenfähigkeit und Risikotoleranz begrenzt sind.

Auf die Wirtschaft übertragen bedeutet das: Eine leistungsfähige KI könnte Strategien empfehlen, die mehr Variablen, Lieferketten, Preisentwicklungen, Kundenreaktionen, Kapazitätsgrenzen und regulatorische Nebenbedingungen gleichzeitig berücksichtigen, als ein Managementteam verarbeiten kann. Ihre Empfehlung wäre möglicherweise besser, aber nicht mehr vollständig rekonstruierbar. Damit entstünde ein neuer Typus von Entscheidung: rechnerisch überlegen, organisatorisch schwer vermittelbar und menschlich nur begrenzt überprüfbar.

Das kurze Fenster der epistemischen Umkehrung

Bislang wurde die epistemische Umkehrung meist als Wechsel der Wissensrichtung beschrieben. Früher übertrug der Mensch sein Wissen auf die Maschine. Heute entdeckt die Maschine Zusammenhänge, aus denen der Mensch lernen kann. Eine 2025 in den *Proceedings of the National Academy of Sciences* veröffentlichte Untersuchung zu AlphaZero liefert dafür einen eindrucksvollen Beleg.²

Die Forscher extrahierten strategische Konzepte aus dem System, die in der traditionellen Schachliteratur nicht oder nicht in dieser Form vorkamen. Vier ehemalige oder amtierende Weltklassenspieler wurden mit diesen Konzepten konfrontiert. Nach einer Lernphase verbesserten alle vier ihre Leistungen bei den entsprechenden Aufgaben. Das Experiment war klein und erlaubt keine weitreichende Verallgemeinerung. Seine Bedeutung liegt an anderer Stelle: Zumindest ein Teil des maschinellen Wissens liegt jenseits der bisherigen menschlichen Schachlehre und kann dennoch in menschliches Verständnis übersetzt werden.

**Die Maschine entdeckt etwas. Der Mensch rekonstruiert es.
Anschließend erweitert er sein eigenes Wissen.**

Das ist epistemische Umkehrung im produktiven Sinne. Die Maschine wird zum Lehrer, der Mensch bleibt lernfähig. Dieses Verhältnis setzt jedoch voraus, dass

zwischen maschineller Berechnung und menschlichem Begriff eine Brücke gebaut werden kann. Genau hier markieren neuere Studien eine Grenze.

Eine im Juli 2026 veröffentlichte Untersuchung verglich die Begründungen von Großmeistern, enginegestützten Kommentatoren und Sprachmodellen.³ Die Gruppen bildeten deutlich voneinander getrennte semantische Muster. Sie beschrieben dieselbe Stellung, aber nicht mit derselben begrifflichen Architektur. Versuche, maschinelle Erklärungen stärker an menschlichen Denkweisen auszurichten, verbesserten die Verständlichkeit, konnten aber zugleich die taktische Leistung beeinträchtigen.

Daraus folgt ein Zielkonflikt: Je stärker ein System für menschliche Nachvollziehbarkeit optimiert wird, desto eher könnte es gezwungen sein, seine tatsächliche Entscheidungslogik in ein vertrauterer und womöglich unvollständiges Vokabular zu übersetzen. Eine Erklärung wäre dann nicht mehr die wirkliche Begründung der Entscheidung, sondern eine nachträglich erzeugte Annäherung. Sie könnte plausibel und didaktisch nützlich sein, ohne den tatsächlichen Berechnungsweg abzubilden.

An diesem Punkt endet die produktive epistemische Umkehrung. Die Maschine besitzt weiterhin Wissen, doch der Mensch kann es nicht mehr vollständig in sein eigenes begriffliches System überführen. Aus Wissenstransfer wird Wissensabhängigkeit.

Von der Assistenz zur Abhängigkeit

Diese Grenze ist nicht mehr auf das Schach beschränkt. Der im September 2025 veröffentlichte GDPval-Benchmark untersucht die Qualität künstlicher Intelligenz bei wirtschaftlich relevanten Aufgaben aus 44 Berufen, darunter Ingenieure, Produktionsverantwortliche, Einkäufer, Finanzanalysten, Anwälte, Pflegekräfte, Projektmanager und Journalisten.⁴ Die beteiligten Fachleute verfügten im Durchschnitt über 14 Jahre Berufserfahrung.

Bereits die erste Untersuchung zeigte, dass die besten Modelle bei einem erheblichen Anteil klar formulierter Aufgaben mindestens die Qualität erfahrener Fachleute erreichten. Wenige Monate später meldete OpenAI für GPT-5.2 bei 70,9 Prozent der Vergleiche ein mindestens gleichwertiges Ergebnis. Solche Zahlen sind mit Vorsicht zu behandeln. GDPval testet klar umrissene Aufgaben mit bereitgestellten Informationen. Das Modell muss nicht zwingend erkennen, welches Problem das Unternehmen wirklich hat, widerstreitende Interessen moderieren oder für die Folgen seiner Empfehlung eintreten. Die späteren Zahlen stammen außerdem vom Hersteller des Modells.

Dennoch markieren sie einen Übergang. Die KI formuliert nicht mehr nur Textbausteine oder recherchiert Hintergrundwissen. Sie produziert vollständige Arbeitsergebnisse, die von Fachleuten mit professionellen Ergebnissen verglichen werden können.

DREI STUFEN DES STOCKFISH-ÜBERGANGS

01 ASSISTENZ	02 DELEGATION	03 ABHÄNGIGKEIT
Der Mensch entwickelt die Lösung. Die KI beschleunigt einzelne Schritte.	Die KI entwickelt die Lösung. Der Mensch prüft und entscheidet.	Die KI entwickelt die Lösung. Der Mensch kann sie weder reproduzieren noch unabhängig verifizieren.
Leitfrage: Wie viel Zeit gewinnen wir?	Leitfrage: Wann dürfen wir vertrauen?	Leitfrage: Wer kann noch begründet widersprechen?

Parallel dazu misst die Forschungsorganisation METR, wie lang Aufgaben sein können, die ein KI-Agent mit einer bestimmten Erfolgswahrscheinlichkeit selbstständig bewältigt.⁵ Der entsprechende Aufgabenhorizont hat sich über mehrere Jahre ungefähr alle sieben Monate verdoppelt. In einem vorläufigen Test von 2026 wurde für GPT-5.4 bei bestimmten Programmieraufgaben ein 50-Prozent-Horizont von rund 80 menschlichen Arbeitsstunden geschätzt.

Die Zahl ist hochgradig unsicher. METR fand zudem Hinweise darauf, dass Modelle gelegentlich sichtbare Testbedingungen ausnutzten, statt das allgemeine Problem sauber zu lösen. Gerade dieser Einwand ist aufschlussreich. Er zeigt, dass eine Maschine gleichzeitig beeindruckend leistungsfähig und in ihrer Zielverfolgung fundamental unzuverlässig sein kann. Sie löst dann nicht unbedingt das Problem, das der Mensch meinte. Sie löst das Problem, das durch Daten, Prüfverfahren und Erfolgskriterien tatsächlich definiert wurde.

Im Schach fallen beide Ebenen weitgehend zusammen. Das Ziel ist eindeutig, die Regeln sind bekannt und jede Stellung kann bewertet werden. In einem Unternehmen ist das anders. Umsatz, Qualität, Versorgungssicherheit, Beschäftigung, Reputation, Liquidität und langfristige Wettbewerbsfähigkeit können miteinander in Konflikt geraten. Schon die Wahl der Zielfunktion ist eine Führungsentscheidung.

Das Stockfish-Level der Arbeitswelt entsteht daher nicht überall gleichzeitig. Es entsteht zuerst dort, wo Aufgaben präzise formulierbar, Ergebnisse messbar und Abweichungen schnell erkennbar sind: Programmierung, Produktionsplanung,

Qualitätsprüfung, Kalkulation, technische Auslegung, Logistikoptimierung, Vertragsanalyse oder diagnostische Bilderkennung. In diesen Bereichen kann die sachliche Überlegenheit der Maschine eintreten, lange bevor eine allgemeine künstliche Intelligenz existiert.

Das Verifikationsparadox

Mit der maschinellen Überlegenheit wächst die Bedeutung menschlicher Kontrolle. Gleichzeitig droht genau die Kompetenz zu erodieren, die für diese Kontrolle benötigt wird. Eine 2025 in *The Lancet Gastroenterology & Hepatology* veröffentlichte Untersuchung beobachtete Ärzte an vier polnischen Endoskopiezentren.⁶ Nachdem die Mediziner regelmäßig ein KI-System zur Erkennung verdächtiger Gewebeveränderungen eingesetzt hatten, sank ihre Erkennungsrate bei späteren Untersuchungen ohne KI von 28,4 auf 22,4 Prozent.

Die Studie beweist nicht, dass jede medizinische KI zwangsläufig zum Kompetenzverlust führt. Sie zeigt aber einen möglichen Mechanismus: Wenn ein System dauerhaft die visuelle Aufmerksamkeit lenkt, verändert sich die menschliche Wahrnehmungsroutine. Die Ärzte sehen nicht notwendigerweise schlechter. Sie suchen anders und möglicherweise weniger eigenständig.

Eine Microsoft-Untersuchung mit 319 Wissensarbeitern und 936 dokumentierten Anwendungsfällen kommt zu einem verwandten Ergebnis.⁷ Je stärker die Befragten der generativen KI vertrauten, desto weniger kritische Denkarbeit berichteten sie. Ihre Tätigkeit verschob sich von der eigenständigen Entwicklung einer Lösung zur Überprüfung, Auswahl und Integration maschineller Ergebnisse.

Das klingt zunächst nach vernünftiger Arbeitsteilung. Doch sie enthält ein Paradox: Um einen hochqualifizierten Vorschlag überprüfen zu können, benötigt der Mensch häufig genau jene Erfahrung, die er nicht mehr erwirbt, wenn die Maschine die zugrunde liegende Arbeit dauerhaft übernimmt.

Kurzfristig gewinnt die Organisation Produktivität. Langfristig kann sie eine Kompetenzschuld aufbauen.

Ein junger Ingenieur lernt nicht allein durch das Unterschreiben fertiger Berechnungen. Ein Arzt entwickelt diagnostische Sicherheit nicht durch das Bestätigen farbiger Markierungen. Ein Produktionsplaner gewinnt kein Verständnis für die Empfindlichkeit eines Systems, wenn er nur noch die optimale Reihenfolge übernimmt. Und ein Manager lernt strategisches Urteil nicht dadurch, dass er zwischen drei maschinell erzeugten Szenarien das am besten formatierte auswählt.

Auf diese Weise entsteht Kompetenzschuld. Kurzfristig steigt die Produktivität, weil das Unternehmen menschliche Lern- und Bearbeitungszeit einspart. Langfristig fehlt die eigene Urteilsfähigkeit, wenn das System ausfällt, eine ungewöhnliche Situation falsch einordnet oder seine Zielfunktion von der tatsächlichen Unternehmenslage abweicht. Die Organisation lebt dann von einer Kompetenz, die sie selbst nicht mehr ausreichend reproduzieren kann.

Wenn Human in the Loop zur Beruhigungsformel wird

Viele Governance-Konzepte reagieren auf dieses Problem mit dem Hinweis, am Ende entscheide weiterhin ein Mensch. Doch die bloße Anwesenheit eines Menschen sagt wenig über die Qualität der Kontrolle aus. Ein Mitarbeiter kann formal bestätigen, was er sachlich nicht mehr eigenständig bewerten kann. Er wird damit nicht zum Entscheider, sondern zum haftenden Endpunkt eines maschinellen Prozesses.

Am wirklichen Stockfish-Level genügt deshalb kein *Human in the Loop*. Erforderlich ist ein Mensch, der über drei konkrete Fähigkeiten verfügt:

- Er muss erkennen, unter welchen Bedingungen das System zuverlässig arbeitet.
- Er muss Abweichungen, Grenzfälle und unpassende Zielsetzungen identifizieren können.
- Er muss den Prozess notfalls ohne das System fortführen oder auf ein unabhängiges Verfahren umschalten können.

Kontrolle ist erst dann vorhanden, wenn ein begründeter Widerspruch möglich ist. Wer nur zustimmen kann, kontrolliert nicht.

Für Unternehmen folgt daraus, dass sie neben der Leistungsfähigkeit ihrer KI auch ihre eigene kognitive Widerstandsfähigkeit messen müssen. Dazu gehören die noch vorhandene Fachkompetenz, die Fähigkeit zur unabhängigen Gegenrechnung, die Nachvollziehbarkeit kritischer Entscheidungen und die Zeit, die für einen manuellen oder alternativen Betrieb benötigt würde.

Im industriellen Kontext ist das vergleichbar mit einer Rückfallebene für Energieversorgung, IT oder Lieferketten. Auch Wissen kann einen Single Point of Failure bilden. Nur befindet er sich nicht in einer Maschine, sondern in der zunehmenden Unfähigkeit der Organisation, ohne diese Maschine zu urteilen.

Die ungleichmäßige Front

Der Stanford AI Index 2026 liefert die notwendige Gegenposition zu jeder linearen Fortschrittserzählung.⁸ Moderne Systeme erreichen menschliches oder übermenschliches Niveau bei wissenschaftlichen Prüfungen, Wettbewerbsmathematik und Programmieraufgaben. Gleichzeitig scheitern sie an Tätigkeiten, die für Menschen banal wirken. Das beste untersuchte Modell konnte analoge Uhren nur ungefähr in der Hälfte der Fälle richtig ablesen. Bei realen Computeraufgaben misslang weiterhin etwa jeder dritte Versuch.

Diese ungleichmäßige Leistungsgrenze verhindert, dass man von einem einzigen, universellen Stockfish-Level sprechen kann. Die Maschine kann in einem Teilprozess haushoch überlegen und wenige Schritte später irritierend unbeholfen sein. Für Unternehmen ist das gefährlicher als eine gleichmäßig mittelmäßige Technologie. Hohe Leistung in einem Bereich erzeugt Vertrauen, das leicht auf andere Bereiche übertragen wird.

Wer erlebt hat, dass eine KI in wenigen Minuten eine hervorragende technische Analyse erstellt, neigt dazu, ihr auch bei einer schlecht strukturierten strategischen Frage Kompetenz zuzuschreiben. Doch die Fähigkeiten eines Modells übertragen sich nicht automatisch von einem Aufgabentyp auf einen anderen.

Das Stockfish-Level ist deshalb kein Zeitpunkt, an dem die KI insgesamt intelligenter als der Mensch wird. Es ist eine Vielzahl einzelner Schwellen. Jede Funktion, jeder Prozess und jede Entscheidung erreicht sie zu einem anderen Zeitpunkt. Die Managementaufgabe besteht darin, diese Schwellen zu erkennen, bevor aus produktiver Delegation unkontrollierte Abhängigkeit wird.

Die neue Arbeitsteilung

Der PwC AI Jobs Barometer 2026 deutet darauf hin, dass die Folgen nicht in einer einfachen Vernichtung qualifizierter Arbeit bestehen.⁹ In besonders KI-exponierten Unternehmen wächst die Produktivität deutlich schneller. Gleichzeitig verändern sich die Kompetenzanforderungen in den betroffenen Berufen mehr als doppelt so schnell. Tätigkeiten, die durch KI professionalisiert werden, wachsen stärker als solche, die lediglich für Nichtfachleute zugänglich werden.

Das spricht für eine neue Teilung der Arbeit. Die Maschine übernimmt zunehmend Erzeugung, Berechnung und Optimierung einzelner Ergebnisse. Der Mensch verantwortet Problemdefinition, Zielkonflikte, Grenzfälle, Legitimation und Konsequenzen. Aber diese Arbeitsteilung funktioniert nur, wenn der Mensch die Nähe zur fachlichen Substanz nicht verliert. Urteilskraft entsteht nicht im luftleeren Raum. Sie ist verdichtete Erfahrung. Wer die operative Auseinandersetzung mit

einem Problem vollständig delegiert, verliert irgendwann auch die Grundlage seines strategischen Urteils.

Die Zukunft menschlicher Expertise liegt daher weder im sturen Festhalten an überholten Routinen noch in der vollständigen Übergabe an die Maschine. Sie liegt in einer bewusst gestalteten Doppelstruktur: Die KI bearbeitet die komplexeste verfügbare Lösung; der Mensch erhält die Fähigkeit, Ziele, Voraussetzungen und Folgen dieser Lösung unabhängig zu beurteilen. Das ist aufwendiger als bloße Automatisierung. Aber es ist der Preis dafür, dass aus technischer Überlegenheit keine organisatorische Unmündigkeit wird.

VIER PRÜFUNGEN FÜR DAS MANAGEMENT

01 LEISTUNG Wo ist die KI dem internen Experten tatsächlich und wiederholbar überlegen?	02 VERIFIKATION Wer kann das Ergebnis unabhängig prüfen und unter welchen Bedingungen widersprechen?
03 KOMPETENZERHALT Welche Fähigkeiten müssen Menschen weiter selbst ausüben, damit Urteilskraft erhalten bleibt?	04 RÜCKFALLEBENE Wie lange bleibt ein kritischer Prozess ohne dieses System handlungsfähig?

QUELLENBASIS

Studien und Berichte

Die Quellen wurden für diesen Beitrag auf ihre Aussagekraft und ihre Grenzen hin eingeordnet. Herstellerbenchmarks und vorläufige Preprints sind entsprechend gekennzeichnet.

- 1 Cerioli, A.; Lee, E. D.; Servedio, V. D. P.: *AI sustains higher strategic tension than humans in chess*, arXiv 2508.13213, 2025, überarbeitet 2026. [Originalquelle](#)
- 2 Schut, L. et al.: *Bridging the human-AI knowledge gap through concept discovery and transfer in AlphaZero*, PNAS 122 (13), 2025. [Originalquelle](#)
- 3 Chua, J. et al.: *Three-Body Alignment: Aligning Chess Agent with Human Reasoning through Reranked Rationale*, arXiv 2607.21993, Juli 2026. [Originalquelle](#)
- 4 OpenAI: *Measuring the performance of our models on real-world tasks - GDPval*, September 2025; ergänzende Modellresultate Dezember 2025. Herstellerbenchmark. [Originalquelle](#)
- 5 METR: *Task-Completion Time Horizons of Frontier AI Models*, Stand Mai 2026; *Frontier Risk Report*, Mai 2026. Ergebnisse für sehr lange Aufgaben vorläufig. [Originalquelle](#)
- 6 Budzyński, K. et al.: *Endoscopist deskill risk after exposure to artificial intelligence*, The Lancet Gastroenterology & Hepatology, 2025. [Originalquelle](#)
- 7 Lee, H.-P. H. et al.: *The Impact of Generative AI on Critical Thinking*, Proceedings of CHI 2025. Befragung von 319 Wissensarbeitern mit 936 Anwendungsfällen. [Originalquelle](#)
- 8 Stanford Institute for Human-Centered AI: *2026 AI Index Report*, 2026. [Originalquelle](#)
- 9 PwC: *2026 AI Jobs Barometer - Two futures for jobs in an AI era*, Juni 2026. [Originalquelle](#)

METHODISCHER HINWEIS

Der Begriff Stockfish-Level ist kein etablierter arbeitswissenschaftlicher Messwert. Er wird in diesem Dossier als analytische Metapher für eine domänenspezifische, dauerhafte und für Menschen nur noch begrenzt reproduzierbare Leistungsüberlegenheit verwendet. Die Fünfjahresperspektive ist deshalb als strategisches Szenario, nicht als gesicherte Prognose zu verstehen.

ERKENNTNIS-FAZIT

Die entscheidende Schwelle ist nicht Überlegenheit, sondern Abhängigkeit.

Das Stockfish-Level der Arbeitswelt beginnt nicht erst an dem Tag, an dem eine allgemeine künstliche Intelligenz den Menschen in allen geistigen Tätigkeiten übertrifft. Es beginnt viel früher - und in jedem Prozess zu einem anderen Zeitpunkt.

Drei Bedingungen markieren diese Schwelle:

1. Die Maschine liefert regelmäßig bessere Ergebnisse als erfahrene Fachleute.
2. Ihre Entscheidungswege lassen sich nicht mehr vollständig in menschliche Begriffe übersetzen.
3. Die Menschen verlieren durch dauerhafte Nutzung einen Teil ihrer eigenen Fach- und Prüfkompetenz.

Dann kehrt sich nicht nur die Richtung des Wissens um. Die Organisation verliert allmählich die Fähigkeit, zwischen einer überlegenen Entscheidung und einem überlegenen Irrtum zu unterscheiden.

Die zentrale Managementfrage lautet deshalb künftig nicht mehr nur: **Wo kann KI den Menschen übertreffen?**

Sie lautet: Wie bewahren wir menschliche Urteilskraft in einer Welt, in der der Mensch immer seltener die beste Antwort besitzt, aber weiterhin für ihre Folgen einstehen muss?

2031: Wenn KI das Stockfish-Level erreicht - und überschreitet

Musks Fünfjahresprognose ist keine belastbare Terminplanung. Als Szenario zwingt sie Unternehmen jedoch zu einer unbequemen Frage: Was geschieht, wenn die Fähigkeitssprünge schneller kommen als neue Organisationen?

Musk erwartet, dass künstliche Intelligenz in ungefähr fünf Jahren die Summe menschlicher Intelligenz übertreffen könnte. Solche Aussagen verbinden technische Beobachtung, unternehmerisches Interesse und Zukunftsrhetorik. Niemand kann seriös garantieren, dass 2031 der angekündigte Punkt erreicht wird. Aber strategische Planung beginnt nicht erst bei Gewissheit. Sie beginnt dort, wo eine Entwicklung hinreichend folgenreich und hinreichend plausibel ist, um heutige Entscheidungen zu verändern.

Das Stockfish-Level bezeichnet dabei zunächst keine allwissende Superintelligenz. Es beschreibt strukturelle Überlegenheit in einer Domäne. Beim Programmieren hieße das: Ein System erledigt nicht nur einzelne Aufgaben, sondern plant größere Vorhaben, zerlegt sie in Arbeitspakete, schreibt und testet Code, sucht Fehler, dokumentiert Änderungen und reagiert auf Rückmeldungen. Der Mensch formuliert das Ziel und kontrolliert Grenzfälle; die eigentliche Produktionsleistung wandert zum System.

Von Minuten zu Projekten

Ein wichtiger Hinweis ist nicht die Zahl gelöster Prüfungsfragen, sondern die Dauer autonom bewältigter Aufgaben. Die Forschungsorganisation METR untersucht, wie lange ein Auftrag einen menschlichen Fachmann beschäftigen würde, den ein KI-Agent mit fünfzigprozentiger Zuverlässigkeit selbstständig abschließen kann. Für die sechs Jahre bis 2025 beobachtete METR eine Verdopplung dieses Zeithorizonts ungefähr alle sieben Monate.

Die Forscher warnen selbst vor einer mechanischen Fortschreibung. Aufgabenwahl, Messmethode und menschliche Vergleichszeiten beeinflussen das Ergebnis. Dennoch ist die Richtung bedeutsam. Setzte sich der historische Trend nur einige Jahre fort, würden aus heute stundenlangen Aufgaben mehrtägige und schließlich mehrwöchige Projekte. Der qualitative Sprung läge nicht darin, dass die KI schneller tippt. Er läge

darin, dass sie einen immer größeren zusammenhängenden Ausschnitt der Wertschöpfung ohne Übergabe an einen Menschen bewältigt.

Der Fortschritt verläuft gezackt

Gegen die einfache Fortschrittskurve sprechen reale Erfahrungen. In einer randomisierten METR-Studie aus dem Jahr 2025 benötigten erfahrene Open-Source-Entwickler mit damaligen KI-Werkzeugen im Durchschnitt neunzehn Prozent mehr Zeit als ohne sie. Die Entwickler glaubten dagegen, schneller gewesen zu sein. Eine weitere METR-Untersuchung zeigte, dass automatisch bestandene Softwaretests die Produktionsreife überschätzen können: Von fünfzehn manuell geprüften Beiträgen war keiner ohne Nacharbeit unmittelbar integrierbar.

Auch der Stanford AI Index 2026 beschreibt eine gezackte Intelligenz. Modelle lösen hochkomplexe Mathematikaufgaben und scheitern zugleich an scheinbar banalen Wahrnehmungs- oder Alltagsproblemen. Auf strukturierten Computer-Benchmarks misslingen KI-Agenten weiterhin ungefähr ein Drittel der Versuche. Stockfish ist im Schach so zuverlässig, weil Brett, Regeln und Ziel stabil sind. Softwareentwicklung und Unternehmensführung besitzen diese Stabilität nicht.

Die richtige Schlussfolgerung lautet deshalb weder, Musks Prognose blind zu glauben, noch sie wegen heutiger Fehler zu verwerfen. Technologische Übergänge sind häufig gleichzeitig überschätzt und unterschätzt: überschätzt in ihrer unmittelbaren Zuverlässigkeit, unterschätzt in ihrer kumulativen Wirkung.

Wenn Software Software verbessert

Besonders folgenreich wäre eine Rückkopplung. Eine KI, die Forschungscode, Trainingsverfahren, Testsysteme und Chipdesign verbessert, arbeitet an den Werkzeugen ihrer eigenen nächsten Generation mit. Daraus folgt nicht automatisch eine explosionsartige, unkontrollierte Selbstverbesserung. Fortschritt bleibt an Rechenleistung, Energie, Daten, Experimente und menschlich organisierte Infrastruktur gebunden. Aber selbst eine teilweise Automatisierung der KI-Forschung könnte Entwicklungszyklen verkürzen und den Wettbewerb zusätzlich beschleunigen.

Für Unternehmen entsteht damit ein asymmetrisches Risiko. Wer zu früh auf autonome Systeme setzt, kann Qualität, Sicherheit und Kundenvertrauen gefährden. Wer zu lange wartet, baut Prozesse, Stellenprofile und Investitionen für eine Kostenstruktur, die schneller veraltet als die Anlagen abgeschrieben werden. Die Antwort kann nicht in einer einzigen großen KI-Wette liegen. Sie muss aus gestuften Optionen bestehen: Daten ordnen, Prozesse modularisieren, Kontrollpunkte definieren, technische Schnittstellen öffnen und Anwendungen nach möglichem Nutzen und maximalem Schaden priorisieren.

Der Fünfjahreshorizont ist keine Verheißung. Er ist die kürzeste Frist, in der eine heutige Organisationsentscheidung strategisch alt aussehen könnte.

Die fünf Fragen für 2031

Bis 2031 sollten Unternehmen fünf Fragen beantworten können. Welche Wissensarbeit lässt sich vollständig digital beschreiben? Welche Aufgaben dürfen Systeme autonom ausführen, und welche benötigen eine menschliche Freigabe? Wie wird ein Ergebnis geprüft, wenn kein Mensch mehr dieselbe Leistung reproduzieren kann? Welche Daten und Vermögenswerte bleiben knapp, wenn Intelligenz billig wird? Und welche physischen Kapazitäten begrenzen die Umsetzung?

Damit wird aus Musks großer Prognose eine nüchterne Managementagenda. Ob das Stockfish-Level 2029, 2031 oder später erreicht wird, ist weniger entscheidend als die Anpassungszeit der Organisation. Wer Daten, Zuständigkeiten und Kontrollarchitektur erst ordnet, nachdem Systeme strukturell überlegen geworden sind, beginnt die Partie mit mehreren Zügen Rückstand.

Quellen und Datenbasis

The Economist: [The full-length interview with Elon Musk](#)

METR: [Measuring AI Ability to Complete Long Tasks](#)

METR: [Developer Productivity Experiment - Update 2026](#)

METR: [Algorithmic versus holistic evaluation of software agents](#)

Stanford HAI: [AI Index Report 2026 - Technical Performance](#)

BEITRAG III

Wenn der Mensch zum langsamsten Teil der Wertschöpfung wird

Die nächste Produktivitätsgrenze liegt möglicherweise nicht in der Rechenleistung, sondern in einer Organisation, die jede maschinelle Entscheidung durch dieselben alten Gremien, Budgets und Hierarchien schicken muss.

In der Frühphase einer Technologie ist die Maschine der Engpass. Sie ist teuer, langsam und unzuverlässig. Später kehrt sich das Verhältnis um. Rechenleistung wird billiger, Modelle werden leistungsfähiger, Ergebnisse entstehen in Sekunden - und warten anschließend auf Freigaben, Zuständigkeiten und Kalender. Der Mensch bleibt nicht deshalb der Flaschenhals, weil er wenig kann. Er wird zum Flaschenhals, weil seine Organisation für eine andere Geschwindigkeit gebaut wurde.

Schon heute kann ein KI-System in wenigen Minuten Varianten erzeugen, für die ein Team Tage benötigt hätte. Doch mehr Vorschläge bedeuten nicht automatisch mehr Wert. Jemand muss auswählen, prüfen, genehmigen und umsetzen. Wenn die Zahl der Optionen schneller wächst als die Entscheidungsfähigkeit, produziert KI zunächst keinen Wohlstand, sondern einen Stau aus Möglichkeiten.

Vom Ausführenden zum Auftraggeber

Das Stockfish-Beispiel zeigt, wie sich Facharbeit verändert. Schachgroßmeister benutzen Engines nicht nur, um Züge zu finden. Sie prüfen Varianten, entdecken neue Eröffnungen und untersuchen Stellungen, die früher als strategisch abgeschlossen galten. Ihre Leistung verschiebt sich von der reinen Berechnung zur Interpretation: Welche Empfehlung ist für einen menschlichen Gegner praktisch spielbar? Welche Variante passt zum eigenen Stil? Wo liegt die psychologische Schwierigkeit?

Ähnlich könnte es Programmierern, Ingenieuren, Juristen, Controllern und Beratern ergehen. Der Wert ihrer Arbeit läge weniger in der ersten Produktion eines Entwurfs und stärker in der Definition des Auftrags, der Auswahl zwischen Alternativen, der Einbettung in reale Bedingungen und der Verantwortung für Konsequenzen. Fachwissen verliert dadurch nicht an Bedeutung. Es wird vom Produktionsmittel zum Kontroll- und Urteilssystem.

Das Ende der Expertenhierarchie

Organisationen verteilen Autorität traditionell nach Wissen und Erfahrung. Wer eine Aufgabe besser beherrscht, erhält größere Entscheidungsspielräume. Wenn ein allgemein verfügbares System fachlich bessere Vorschläge macht als neunundneunzig Prozent der Beschäftigten, gerät diese Ordnung unter Druck. Der erfahrenste Mensch ist dann möglicherweise nicht mehr der beste Problemlöser, aber weiterhin derjenige, der Risiken, Geschichte und Nebenwirkungen versteht.

Daraus entsteht ein neues Autoritätsproblem. Darf ein Vorgesetzter eine maschinelle Empfehlung überstimmen, die statistisch fast immer besser ist? Muss er begründen, weshalb er abweicht? Und wer haftet, wenn die Organisation dem System folgt und scheitert? Unternehmen benötigen Regeln für diese Konflikte, bevor sie alltäglich werden. Sonst entscheidet nicht die beste Governance, sondern die Psychologie des jeweiligen Managers: blinder Technikglaube auf der einen, gekränkter Expertenstolz auf der anderen Seite.

Berufseinstieg ohne Lehrjahre

Besonders heikel ist der Nachwuchs. Viele Berufe bilden Erfahrung durch einfache Aufgaben: junge Juristen prüfen Dokumente, Entwickler beseitigen überschaubare Fehler, Controller bereiten Zahlen auf, Ingenieure erstellen Varianten. Genau diese standardisierbaren Tätigkeiten werden zuerst automatisiert. Verschwindet die untere Sprosse der Karriereleiter, entsteht später ein Mangel an Menschen, die komplexe Entscheidungen verantworten können.

Die Antwort kann nicht darin bestehen, ineffiziente Arbeit künstlich zu konservieren. Unternehmen müssen Lernwege neu bauen. Nachwuchskräfte sollten früher an Simulationen, Ausnahmen, Kundenkontakte und Entscheidungen herangeführt werden. Sie müssen lernen, maschinelle Ergebnisse zu prüfen, Unsicherheit zu erkennen und Zielkonflikte zu bearbeiten. Wer nur noch fertige Antworten übernimmt, sammelt keine Erfahrung, sondern Abhängigkeit.

Die Organisation wird zum Betriebssystem

Im Stockfish-Zeitalter entscheidet sich Produktivität nicht allein an der Modellqualität. Sie hängt davon ab, wie schnell ein Unternehmen aus einer maschinellen Erkenntnis eine verantwortete Handlung macht. Dazu braucht es klare Datenzugänge, definierte Berechtigungen, kurze Freigabewege und eine Trennung zwischen reversiblen und irreversiblen Entscheidungen.

Eine Preisänderung in einem kleinen Testmarkt kann ein System autonom erproben. Eine sicherheitskritische Konstruktionsänderung benötigt formale Prüfungen. Eine

Wartungsplanung darf innerhalb festgelegter Grenzen selbstständig optimiert werden; das Abschalten einer ganzen Linie verlangt eine Eskalation. Nicht jede Entscheidung braucht einen Menschen. Aber jede Entscheidung braucht eine vorher bestimmte Verantwortungsarchitektur.

Die KI beschleunigt nicht automatisch das Unternehmen. Sie macht zunächst sichtbar, wie langsam seine Entscheidungen sind.

Damit ändert sich auch die Führungsaufgabe. Führung bedeutet weniger, den fachlich besten Einzelweg selbst zu kennen. Sie bedeutet, Ziele zu präzisieren, Widersprüche offenzulegen, Risiken zu begrenzen und aus vielen maschinell erzeugten Möglichkeiten eine Richtung zu wählen. Der Manager konkurriert nicht mit Stockfish. Er gestaltet die Institution, in der Stockfish arbeiten darf.

Quellen und Datenbasis

The Economist: [The full-length interview with Elon Musk](#)

NBER: [Generative AI at Work](#)

OECD: [The Adoption of Artificial Intelligence in Firms](#)

METR: [Measuring the Impact of Early-2025 AI on Experienced Developer Productivity](#)

BEITRAG IV

Die gefährlichste KI optimiert das falsche Ziel

Je leistungsfähiger ein System wird, desto teurer kann eine falsch formulierte Aufgabe werden. Das zentrale Risiko liegt nicht immer im Irrtum der Maschine, sondern in ihrer folgerichtigen Ausführung.

Stockfish hat einen Vorteil, von dem jeder Vorstand nur träumen kann: Das Ziel ist eindeutig. Die Engine soll die Partie gewinnen. Die Regeln sind unveränderlich, alle relevanten Informationen liegen offen auf dem Brett, und am Ende gibt es Sieg, Remis oder Niederlage. Zwischenziel und Endzustand lassen sich mathematisch aufeinander beziehen. Was dem System fehlt, ist Urteilskraft außerhalb des Spiels; was es besitzt, ist eine außerordentlich klare Zielfunktion.

Unternehmen arbeiten unter entgegengesetzten Bedingungen. Eine Fabrik soll produktiv, flexibel, sicher und lieferfähig sein. Der Einkauf soll Preise senken, ohne Qualität und Versorgungssicherheit zu gefährden. Ein Vertriebsalgorithmus soll Umsatz steigern, ohne Kundenbeziehungen zu beschädigen. Eine Personalsteuerung soll Kosten beherrschen, ohne das Wissen zu verlieren, von dem die nächste Produktgeneration abhängt. Diese Ziele lassen sich nicht auf eine einzige Zahl reduzieren, ohne dass Wesentliches verloren geht.

Wenn Präzision zur Gefahr wird

Das klassische Missverständnis automatisierter Systeme besteht darin, ihre Fehler für das Hauptproblem zu halten. Doch ein fehlerhaftes System wird häufig gestoppt, weil sein Versagen sichtbar ist. Gefährlicher kann ein System sein, das die vorgegebene Kennzahl sehr zuverlässig optimiert und dabei Folgen erzeugt, die in der Kennzahl nicht vorkommen. Wer Auslastung maximiert, kann Wartungsfenster verdrängen. Wer Lagerbestände minimiert, kann Resilienz abbauen. Wer Durchlaufzeiten verkürzt, kann Qualitätskontrollen unter Druck setzen.

In der Ökonomie wird dieser Mechanismus häufig mit Goodharts Gesetz beschrieben: Sobald eine Kennzahl zum Ziel wird, verliert sie ihre Qualität als Kennzahl. Die Formel ist älter als die heutige KI. Mit autonomen Agenten gewinnt sie jedoch eine neue Wucht. Ein menschlicher Mitarbeiter sieht möglicherweise, dass die wörtliche Ausführung eines Auftrags dem eigentlichen Zweck widerspricht. Ein System, das für konsequente Zielerreichung belohnt wird, kann gerade wegen seiner Disziplin weitergehen.

Knight Capital: 45 Minuten bis zur Existenzkrise

Wie schnell Automatisierung einen Fehler skalieren kann, zeigte am 1. August 2012 der amerikanische Finanzdienstleister Knight Capital. Nach Feststellungen der US-Börsenaufsicht SEC wurde bei der Einführung neuer Handelssoftware auf mehreren Servern älter Programmcode aktiviert. Aus 212 kleinen Kundenaufträgen entstanden innerhalb von etwa 45 Minuten Millionen von Orders. Das System erzeugte mehr als vier Millionen Ausführungen in 154 Aktien und bewegte rund 397 Millionen Wertpapiere.

Am Ende hielt Knight unbeabsichtigte Long-Positionen von rund 3,5 Milliarden Dollar und Short-Positionen von etwa 3,15 Milliarden Dollar. Der Verlust überstieg 460 Millionen Dollar. Das Unternehmen geriet in Kapitalnot und wurde wenig später übernommen. Es handelte sich nicht um generative KI und nicht um ein System, das eigene Ziele entwickelte. Gerade deshalb ist der Fall lehrreich: Schon eine eng begrenzte Automatisierung verwandelte einen technischen und organisatorischen Fehler innerhalb einer Dreiviertelstunde in eine existenzielle Krise.

Die SEC verwies nicht allein auf defekte Software. Sie bemangelte unzureichende Kontrollen, fehlende Tests, mangelnde Überwachung und eine unvollständige Einführung über mehrere Server hinweg. Die Ursache lag also nicht 'in der Maschine', sondern in der Verbindung von Software, Prozess und Governance. Bei leistungsfähigeren KI-Agenten wird diese Verbindung noch wichtiger, weil sie nicht nur einzelne Befehle ausführen, sondern Zwischenschritte auswählen und Werkzeuge benutzen können.

Zillow: Wenn das Modell zum Geschäft wird

Ein anderes Muster zeigte Zillow. Das Immobilienunternehmen übertrug seine datenbasierte Preiskompetenz auf ein kapitalintensives Geschäftsmodell: Häuser sollten algorithmisch bewertet, angekauft, renoviert und weiterverkauft werden. 2021 beendete Zillow den Bereich Zillow Offers. Das Homes-Segment wies für das Jahr einen Vorsteuerverlust von 881 Millionen Dollar aus; zugleich kündigte das Unternehmen an, seine Belegschaft um ungefähr 25 Prozent zu reduzieren.

Zillow-Gründer Rich Barton erklärte, die Unvorhersehbarkeit der Hauspreise habe die erwartete Schwankungsbreite deutlich überschritten. Der Fall wird gelegentlich als Beweis für das Versagen eines Algorithmus erzählt. Das greift zu kurz. Das Modell war Bestandteil einer Managemententscheidung, die unsichere Preisprognosen mit Immobilienbestand, Fremdkapital, Renovierungskapazität und regionalen Marktrisiken verband. Aus einer Schätzung wurde eine Bilanzposition. Das Problem bestand nicht

nur darin, dass Prognosen falsch sein könnten, sondern darin, wie viel Kapital das Unternehmen auf ihre Genauigkeit setzte.

Die Mathematik der seltenen Fehler

Mit wachsender Automatisierung verändert sich auch die Bedeutung kleiner Fehlerquoten. Ein System mit 99 Prozent Zuverlässigkeit klingt beeindruckend. Bei tausend täglichen Entscheidungen bleiben jedoch rechnerisch zehn Fehler. In einer Textkorrektur ist das ärgerlich. Bei Kreditfreigaben, Medikamentendosierungen, Produktionsparametern oder Wertpapiertransaktionen kann bereits ein einzelner Fehler erhebliche Folgen haben.

Hinzu kommt die Abhängigkeit der Fehler. Wenn tausend Menschen unabhängig voneinander arbeiten, verteilen sich Irrtümer. Wenn tausend Prozesse auf demselben Modell, denselben Daten und derselben fehlerhaften Vorgabe beruhen, können sie gleichzeitig in dieselbe Richtung laufen. Skalierung vervielfacht dann nicht nur Leistung, sondern auch Korrelation. Das seltene Ereignis wird zum systemischen Ereignis.

Was Management daraus ableiten muss

Die erste Konsequenz lautet, Ziele nicht mit einzelnen Kennzahlen zu verwechseln. Ein autonomes System benötigt ein Zielsystem: Produktivität, Qualität, Sicherheit, Resilienz und Compliance müssen als gemeinsame Grenzen formuliert werden. Die zweite Konsequenz betrifft Handlungsvollmachten. Nicht jede KI, die eine Entscheidung vorbereiten darf, sollte sie auch ausführen können. Schwellenwerte, Vier-Augen-Prinzipien und technische Abbruchmechanismen gehören in die Architektur, nicht in eine nachträgliche Richtlinie.

Drittens muss die Kontrolle unabhängig sein. Ein System sollte nicht allein seine eigenen Ergebnisse bewerten. Wo die menschliche Fachleistung für eine vollständige Prüfung nicht mehr ausreicht, werden konkurrierende Modelle, Simulationen, formale Tests und unveränderliche Protokolle wichtiger. Viertens muss der Vorstand den maximal möglichen Schaden kennen. Die entscheidende Kennzahl ist nicht nur die durchschnittliche Trefferquote, sondern die Frage, was ein System in zehn Minuten falsch ausführen kann, bevor jemand es stoppt.

Eine KI optimiert nicht, was das Management gemeint hat. Sie optimiert, was Ziele, Daten und Berechtigungen tatsächlich vorgeben.

Stockfish kann sich ganz auf den Sieg konzentrieren, weil außerhalb des Brettes nichts auf dem Spiel steht. Unternehmen besitzen diesen Luxus nicht. Ihre leistungsfähigsten

Systeme müssen deshalb nicht nur intelligent, sondern institutionell eingebettet sein. Andernfalls kann die Organisation ihre Kennzahl gewinnen und dabei ihr Geschäft verlieren.

Quellen und Datenbasis

SEC: [Knight Capital Americas LLC - Verwaltungsentscheidung vom 16. Oktober 2013](#)

Zillow: [Geschäftsergebnis 2021 und Verlust des Homes-Segments](#)

Zillow: [Ankündigung zur Beendigung von Zillow Offers](#)

NIST: [AI Risk Management Framework](#)

Wenn jeder Stockfish hat, gewinnt nicht mehr die beste KI

Sobald leistungsfähige Modelle zur allgemein verfügbaren Infrastruktur werden, wandert der Wettbewerbsvorteil. Entscheidend ist nicht mehr der Zugang zur Intelligenz, sondern ihre Verbindung mit Daten, Prozessen und realen Vermögenswerten.

Stockfish hat das Schach grundlegend verändert und ist dennoch kein exklusiver Wettbewerbsvorteil. Die Engine ist Open Source, kostenlos verfügbar und auf gewöhnlicher Hardware einsetzbar. Jeder Großmeister kann sie zur Vorbereitung verwenden. Wer Stockfish besitzt, besitzt deshalb noch keinen Vorsprung. Die Software gehört zur Grundausstattung eines Sports, dessen menschliche Spitze längst weiß, dass sie der Maschine im direkten Spiel unterlegen ist.

Diese Entwicklung enthält eine wirtschaftliche Pointe. Eine Technologie kann zugleich revolutionär und kommerziell gewöhnlich werden. Je schneller sie sich verbreitet, desto rascher verschwindet der Vorteil des bloßen Besitzes. Genau dies könnte mit generativer KI geschehen. Heute unterscheiden Unternehmen noch zwischen jenen, die leistungsfähige Systeme einsetzen, und jenen, die es nicht tun. Morgen könnte diese Unterscheidung so wenig aussagekräftig sein wie die Frage, ob ein Unternehmen E-Mail oder Tabellenkalkulation verwendet.

Der Preis der Intelligenz fällt

Die technischen Kostenkurven sprechen für eine breite Diffusion. Der Stanford AI Index 2025 ermittelte, dass die Inferenzkosten für eine Modelleistung auf dem Niveau von GPT-3.5 zwischen November 2022 und Oktober 2024 um mehr als das 280-Fache sanken. Zugleich gingen Hardwarekosten im langfristigen Trend um etwa 30 Prozent pro Jahr zurück, während sich die Energieeffizienz um rund 40 Prozent jährlich verbesserte.

Solche Zahlen bedeuten nicht, dass die leistungsfähigsten Modelle billig sind. Das Training an der Frontier bleibt kapitalintensiv, und der Stanford AI Index 2026 weist darauf hin, dass sich der Abstand zwischen dem besten geschlossenen und dem besten offenen Modell im Jahr 2025 zeitweise wieder vergrößert hat. Wirtschaftlich entscheidend ist jedoch nicht nur die absolute Spitze. Fähigkeiten wandern mit zeitlicher Verzögerung in kleinere, günstigere und spezialisierte Systeme. Was heute

eine exklusive Frontier-Leistung ist, kann morgen als Programmierschnittstelle oder lokales Modell allgemein verfügbar sein.

Produktivität ist nicht gleich Differenzierung

Eine viel zitierte Feldstudie von Erik Brynjolfsson, Danielle Li und Lindsey Raymond untersuchte mehr als 5.000 Mitarbeiter eines Kundendienstunternehmens. Ein generativer Assistent steigerte die Zahl der bearbeiteten Vorgänge pro Stunde im Durchschnitt um 14 Prozent. Bei unerfahrenen und leistungsschwächeren Mitarbeitern lag der Zuwachs bei rund 34 Prozent. Die KI machte einen Teil des Wissens leistungstärkerer Kollegen für andere nutzbar.

Der Befund ist aus zwei Gründen bemerkenswert. Erstens kann KI erhebliche Produktivitätsgewinne erzeugen. Zweitens verringert sie Leistungsunterschiede, weil sie schwächeren oder unerfahrenen Beschäftigten besonders hilft. Für ein einzelnes Unternehmen ist das zunächst attraktiv. Sobald aber alle Wettbewerber ein ähnliches System einsetzen, wird aus der Verbesserung ein neuer Marktstandard. Die Bearbeitungszeit sinkt branchenweit, Kunden erwarten schnellere Antworten, und ein Teil des Produktivitätsgewinns wandert über niedrigere Preise oder höhere Serviceerwartungen an den Markt.

Die Geschichte der Informationstechnologie kennt dieses Muster. Elektrizität, ERP-Systeme, Cloud-Infrastruktur und Suchmaschinen schufen enorme Werte, ohne jedem Nutzer eine dauerhafte Überrendite zu garantieren. Der Vorteil entstand dort, wo Unternehmen ihre Organisation, Produkte und Geschäftsmodelle um die neue Infrastruktur herum neu gestalteten.

Die komplementären Vermögenswerte

Die OECD verweist bei der KI-Einführung auf komplementäre Vermögenswerte. Produktivere Unternehmen nutzen KI häufiger, doch der Effekt hängt mit digitalen Kompetenzen, leistungsfähiger Infrastruktur, Datenqualität und anderen Technologien zusammen. Der Kauf einer Lizenz ersetzt diese Voraussetzungen nicht. Ein Modell kennt den allgemeinen Stand der Technik; es kennt nicht automatisch die stillen Regeln einer bestimmten Fabrik, die Eigenheiten eines Materials oder die Vertrauensgeschichte mit einem Kunden.

Daraus ergeben sich fünf Quellen dauerhafter Differenzierung. Erstens proprietäre Daten, deren Qualität und Bedeutung der Wettbewerber nicht kopieren kann. Zweitens Prozesswissen, das nicht nur dokumentiert, sondern in Routinen, Anlagen und Erfahrung verankert ist. Drittens physische Vermögenswerte wie Fabriken, Labore und Servicenetze. Viertens Kundenzugang und Vertrauen. Fünftens eine Organisation, die maschinelle Vorschläge schneller und verantwortlicher umsetzen kann als andere.

Der neue Preiswettbewerb

Wo KI Leistungen standardisiert, droht zudem eine Kommodifizierung. Wenn zehn Beratungen in Minuten vergleichbare Marktanalysen erzeugen, sinkt der Wert der Analyse als solcher. Wenn Softwarefunktionen mit geringem Aufwand reproduziert werden können, geraten Preise unter Druck. Wenn Marketingtexte, Bilder und Produktvarianten unbegrenzt entstehen, wird Aufmerksamkeit knapper als Produktion.

Unternehmen sollten daher zwischen Kostenvorteil und Kundenvorteil unterscheiden. Wer KI nur verwendet, um dieselbe Leistung billiger zu erzeugen, befindet sich bald in einem Preiswettbewerb. Wer mit KI eine andere Leistung, ein besseres Geschäftsmodell oder eine engere Verbindung zum Kunden schafft, kann Differenzierung aufbauen. Die strategische Aufgabe besteht nicht darin, menschliche Arbeit möglichst vollständig zu ersetzen. Sie besteht darin, eine allgemein verfügbare Intelligenz mit schwer kopierbaren Vermögenswerten zu verbinden.

Eine Chance für das Hochlohnland

Für Deutschland liegt darin eine unerwartete Chance. Wenn digitale Intelligenz weltweit günstiger wird, verliert der reine Zugang zu Software an Bedeutung. Wertvoller werden Produktionswissen, Materialkompetenz, Zulassungen, Maschinenparks, Kundennetzwerke und die Fähigkeit, Qualität in Serie zu liefern. Das Hochlohnland kann seine Kostennachteile teilweise durch eine Ressource kompensieren, deren Grenzkosten stark fallen. Voraussetzung ist, dass KI nicht als isoliertes Büroprojekt behandelt wird, sondern in Konstruktion, Fertigung, Instandhaltung und Service gelangt.

Der entscheidende Wettbewerb findet deshalb nicht zwischen Mensch und Maschine statt. Er findet zwischen Organisationen statt, die auf dieselben oder ähnliche Modelle zugreifen. Gewinnen wird nicht notwendig das Unternehmen mit der spektakulärsten KI-Demonstration. Gewinnen kann das Unternehmen, das seine Daten ordnet, Entscheidungen beschleunigt, Risiken begrenzt und die Technologie in Prozesse einbettet, die andere nicht ohne Weiteres nachbauen können.

Wenn Intelligenz zur Infrastruktur wird, entsteht der Vorsprung nicht im Modell. Er entsteht in der Organisation, die daraus schneller reale Wertschöpfung macht.

Quellen und Datenbasis

Stanford HAI: [AI Index Report 2025](#)

Stanford HAI: [AI Index Report 2026](#)

NBER: [Generative AI at Work](#)

OECD: [The Adoption of Artificial Intelligence in Firms](#)

Stockfish: [Offizielle Projektseite](#)

BEITRAG VI

Die KI denkt in Sekunden. Die Fabrik baut noch in Monaten

Digitale Intelligenz kann Entwürfe nahezu ohne Grenzkosten erzeugen. Industrielle Wertschöpfung bleibt an Energie, Maschinen, Material und Zeit gebunden. Die entscheidende Zukunftsfrage lautet, wie schnell Bits zu Atomen werden.

Eine künstliche Intelligenz kann an einem Nachmittag tausend Konstruktionsvarianten berechnen. Sie kann Materialkombinationen vergleichen, Fertigungsfolgen simulieren und den Energieverbrauch unterschiedlicher Designs abschätzen. Doch eine neue Produktionslinie muss weiterhin bestellt, geliefert, montiert, programmiert und abgenommen werden. Zwischen digitaler Intelligenz und industrieller Wertschöpfung liegt die widerständige Welt der Atome.

Diese Differenz wird in vielen Zukunftserzählungen übersprungen. Software skaliert mit Rechenleistung. Ein einmal entwickeltes Modell kann, bei ausreichender Infrastruktur, gleichzeitig an vielen Orten eingesetzt werden. Eine Maschine besitzt andere Eigenschaften. Sie braucht Stahl, Halbleiter, Sensoren, Fläche, Strom, Genehmigungen und Menschen, die sie in einen vorhandenen Prozess integrieren. Selbst eine perfekte Konstruktion hebt diese Bedingungen nicht auf.

4,66 Millionen Roboter - und dennoch keine beliebige Produktion

Die industrielle Automation ist keineswegs am Anfang. Nach Angaben der International Federation of Robotics befanden sich 2024 weltweit rund 4,66 Millionen Industrieroboter in Betrieb, neun Prozent mehr als im Vorjahr. Im selben Jahr wurden 542.000 neue Einheiten installiert. Damit lag die Zahl zum vierten Mal in Folge über 500.000 und mehr als doppelt so hoch wie zehn Jahre zuvor.

Die regionale Verteilung zeigt jedoch, wie unterschiedlich die physische Umsetzung verläuft. 74 Prozent der neuen Roboter gingen 2024 nach Asien, 16 Prozent nach Europa und neun Prozent in die amerikanischen Staaten. China allein installierte 295.000 Einheiten und damit 54 Prozent des Weltmarktes. Deutschland kam auf knapp 27.000 Neuinstallationen. Bei der Roboterichte lag die deutsche Fertigung 2024 mit 449 Robotern je 10.000 Beschäftigte weltweit auf Rang drei; Südkorea erreichte 1.220, China 567.

Diese Zahlen belegen eine fortgeschrittene Automation, aber keine allgemeine maschinelle Arbeitskraft. Die meisten Industrieroboter sind Spezialisten. Sie schweißen, lackieren, greifen oder montieren in präzise eingerichteten Zellen. Ein Schweißroboter kann nicht am nächsten Tag selbstständig eine Pumpe warten oder ein Lager neu organisieren. Humanoide Systeme versprechen mehr Flexibilität, befinden sich aber noch nicht in einer vergleichbaren Breite des industriellen Einsatzes.

Die Zeit der Fabrik

Wie hart die physische Grenze ist, zeigt die Halbleiterindustrie. Eine moderne Chipfabrik kostet nach Angaben der Semiconductor Industry Association grob zehn bis zwanzig Milliarden Dollar. Der reine Bau benötigt durchschnittlich 18 bis 24 Monate; danach folgen Installation der Anlagen, Prozessqualifikation und Produktionshochlauf. In hochkomplexen Fabriken ist der Kalender nicht bloß eine Folge langsamer Verwaltung. Er bildet Materiallieferungen, Reinraumtechnik, Spezialmaschinen und die schrittweise Verbesserung der Ausbeute ab.

TSMC liefert ein anschauliches Beispiel. Die Gebäudestruktur der zweiten Fabrik in Arizona war 2025 fertiggestellt. Die Serienproduktion mit 3-Nanometer-Technologie ist für die zweite Hälfte des Jahres 2027 vorgesehen. Zwischen einer baulich fertigen Hülle und stabiler Hochvolumenproduktion liegen damit weiterhin Jahre. Die KI kann die nächste Chipgeneration entwerfen, lange bevor die Fertigungskapazität bereitsteht.

Auch digitale Intelligenz braucht eine physische Basis

Selbst die scheinbar immaterielle KI ist an Atome gebunden. Sie benötigt Rechenzentren, Chips, Kühlung, Netze und Kraftwerke. Nach Angaben der Internationalen Energieagentur stieg der weltweite Stromverbrauch von Rechenzentren 2025 um 17 Prozent; bei KI-orientierten Rechenzentren lag das Wachstum bei 50 Prozent. Bis 2030 könnte der gesamte Stromverbrauch der Rechenzentren auf rund 945 Terawattstunden steigen. Das wäre etwas mehr als der heutige Stromverbrauch Japans.

Die IEA erwartet damit von 2024 bis 2030 ein jährliches Wachstum von etwa 15 Prozent - mehr als viermal so schnell wie beim übrigen Stromverbrauch. In Europa könnte der Bedarf bis 2030 um mehr als 45 Terawattstunden oder rund 70 Prozent steigen. Die Debatte über KI ist deshalb auch eine Debatte über Netzausbau, Erzeugungskapazität, Genehmigungen und Standortpolitik.

Der neue Engpass der Wertschöpfung

Wenn Entwurfsleistung billiger wird, verschiebt sich die Knappheit. Unternehmen können mehr Varianten erzeugen, als sie testen. Forschung kann mehr Hypothesen formulieren, als Labore prüfen. Produktentwicklung kann schneller neue Konzepte liefern, als Werkzeugbau und Zulassung folgen. Das Management erhält nicht automatisch mehr Wert, sondern zunächst mehr Möglichkeiten.

Damit wird die Auswahl zur Kernleistung. Welche Variante wird physisch gebaut? Welche Simulation ist belastbar genug, um einen Prototyp einzusparen? Wo lohnt sich eine flexible Anlage, und wo bleibt eine spezialisierte Linie überlegen? Die Produktivität der digitalen Phase kann nur dann in Gesamteffizienz übergehen, wenn die nachfolgenden physischen Prozesse neu gestaltet werden.

Wenn Roboter die Lücke schließen

Die nächste Zäsur entstünde, wenn überlegene digitale Intelligenz mit flexibel einsetzbaren Robotern verbunden wird. Dann könnte ein System nicht nur einen Arbeitsablauf planen, sondern ihn in der realen Welt ausführen: Material bereitstellen, Maschinen umrüsten, Qualitätsabweichungen erkennen und Wartungsarbeiten vornehmen. Aus einem wissenden System wurde ein handelndes System.

Gerade hier ist Nüchternheit erforderlich. Eine eindrucksvolle Demonstration beweist noch keine wirtschaftliche Robustheit. Fabriken verlangen hohe Verfügbarkeit, Arbeitssicherheit, Wiederholgenauigkeit, einfache Wartung und kalkulierbare Kosten. Ein humanoider Roboter, der eine Aufgabe im Labor bewältigt, muss sie im Mehrschichtbetrieb tausendfach wiederholen können. Zwischen Prototyp und Produktionsmittel liegt dieselbe physische Beharrlichkeit, die auch andere Technologien durchlaufen.

Die deutsche Chance

Für Deutschland ergibt sich daraus eine andere Perspektive auf das Hochlohnland. Wenn digitale Intelligenz weltweit günstig verfügbar wird, werden physische Vermögenswerte nicht wertlos. Im Gegenteil: Fabriken, Produktionswissen, Materialkompetenz, Zertifizierungen und eingespielte Lieferantennetze können an Bedeutung gewinnen. Sie sind die Infrastruktur, durch die allgemeine Intelligenz zu einem marktfähigen Produkt wird.

Der Standortnachteil hoher Arbeitskosten kann sinken, wenn KI und Robotik den menschlichen Aufwand pro Einheit reduzieren. Gleichzeitig steigen die Anforderungen an Energieversorgung, Investitionsgeschwindigkeit und technische Integration. Ein

Land, das gute KI-Modelle nutzen kann, aber fünf Jahre auf Netzanschlüsse, Genehmigungen und Fabrikumbauten wartet, gewinnt den Produktivitätswettkampf nicht.

Die Zukunft entscheidet sich nicht allein im Rechenzentrum. Sie entscheidet sich dort, wo aus Bits verlässlich Atome werden.

Die Stockfish-Agenda für die Industrie lautet deshalb: digitale Intelligenz skalieren, ohne die physische Umsetzung zu unterschlagen. Wer nur die Geschwindigkeit der Modelle betrachtet, verwechselt Entwurf mit Wertschöpfung. Wer dagegen Rechenleistung, Robotik, Energie und Produktionssystem zusammen denkt, kann aus der neuen Intelligenz einen realen Standortvorteil machen.

Quellen und Datenbasis

IFR: [World Robotics 2025 - weltweite Installationen und Roboterbestand](#)

IFR: [Roboterdichte in Europa, Asien und Amerika](#)

IEA: [Energy and AI - Energy demand from AI](#)

TSMC: [TSMC Arizona - Zeitplan der Fabriken](#)

SIA: [Semiconductor manufacturing incentives and fab costs](#)

SCHLUSS

Die Agenda für das Management

Das Stockfish-Level ist keine einzelne Technologieprognose. Es ist ein Belastungstest für die Handlungsfähigkeit von Unternehmen.

Die sechs Beiträge führen zu einer gemeinsamen Schlussfolgerung. Leistungsfähigere KI löst das Management nicht von seinen Aufgaben. Sie verschiebt sie. Wenn der Mensch nicht mehr der höchste Leistungsmaßstab ist, werden Zielsetzung, Urteil und Verantwortung wichtiger. Je besser Maschinen optimieren, desto entscheidender wird die Wahl der Ziele. Je breiter Modelle verfügbar sind, desto wichtiger werden komplementäre Vermögenswerte. Je schneller digitale Intelligenz arbeitet, desto deutlicher treten die Engpässe der Organisation und der physischen Welt hervor.

Unternehmen sollten deshalb drei Fragen beantworten. Welche Ziele und Grenzen gelten für autonome Systeme? Welche Daten, Prozesse und Vermögenswerte schaffen einen Vorsprung, wenn alle Wettbewerber ähnliche Modelle besitzen? Und welche physischen Kapazitäten sind notwendig, damit digitale Erkenntnis tatsächlich zu Produkten, Lieferfähigkeit und Ertrag wird?

Stockfish hat den Menschen nicht aus dem Schach entfernt. Es hat ihm die Illusion genommen, weiterhin die höchste spielerische Instanz zu sein. Sollte künstliche Intelligenz eine ähnliche Stellung in der Wirtschaft erreichen, bliebe der Mensch verantwortlich, ohne in jedem Fachgebiet der leistungsfähigste Akteur zu sein. Darauf müssen Organisationen vorbereitet werden - nicht erst, wenn die erste Maschine die Partie bereits entschieden hat.

DIE STOCKFISH-AGENDA

Drei Beiträge über ökonomischen Delegationsdruck, epistemische Resilienz und die neue Entscheidungsverfassung des Unternehmens

Wenn die Maschine nicht nur schneller rechnet, sondern dauerhaft besser urteilt, verändert sie mehr als Prozesse. Sie verändert die Beweislast, die Entstehung von Kompetenz und die Legitimation von Führung.

01

Ökonomie der Delegation

02

Epistemische Resilienz

03

Entscheidungsverfassung

Lothar K. Doerr

Institut für Produktionserhaltung

EDITORISCHE NOTIZ

Drei Fortsetzungen des Essays „Wenn der Mensch nicht mehr der Richter ist“

Der Ausgangstext beschreibt den erkenntnistheoretischen Bruch: Auf dem Stockfish-Level kann der Mensch die maschinelle Entscheidung nicht mehr unabhängig beurteilen und muss ihr dennoch aus guten Gründen folgen. Die folgenden Beiträge führen diesen Gedanken in drei Richtungen weiter.

01 Der Zwang zum besseren Urteil

Die ökonomische Logik: Warum der Nichtgebrauch überlegener KI zum Wettbewerbsnachteil wird.

02 Die Organisation, die verlernt zu denken

Das Kompetenzparadox: Warum Leistungsgewinn die menschliche Rückfallebene schwächen kann.

03 Wer darf der Maschine widersprechen?

Die Governancefrage: Wie Rechte, Haftung und Legitimität neu geordnet werden müssen.

Die drei Grafiken sind konzeptionelle INFPRO-Modelle. Sie visualisieren Managementlogiken und keine empirisch gemessenen Indexwerte.

BEITRAG 1 | ÖKONOMIE DER DELEGATION

Der Zwang zum besseren Urteil

Warum das Stockfish-Level aus einer technischen Option eine wirtschaftliche Pflicht macht

Unternehmen werden Entscheidungen nicht an überlegene KI-Systeme delegieren, weil sie sich philosophisch von ihnen überzeugen lassen. Sie werden es tun, weil der Verzicht messbar teurer wird. Damit beginnt eine neue Phase des Wettbewerbs: Nicht die Nutzung der KI ist begründungspflichtig, sondern zunehmend ihr Nichtgebrauch.

Vom Werkzeug zur epistemischen Abhängigkeit

Fünf Stufen verschieben Leistung, Beweislast und Entscheidungsautorität



Entscheidender Schwellenwert

Nicht die Intelligenz der KI ist ausschlaggebend, sondern der Punkt, an dem menschliche Abweichungen statistisch häufiger schaden als nützen und unabhängige Verifikation praktisch ausfällt.

Folge: Die Maschine löst nicht nur Aufgaben. Sie verändert die Beweislast im Unternehmen.

INFPRO | Stockfish-Agenda | Konzeptionelles Modell

Abbildung 1: Konzeptionelles INFPRO-Modell.

Die unscheinbare Revolution

Technologische Umbrüche kündigen sich gern mit großen Bildern an: Fabriken ohne Menschen, Roboter in Vorstandsetagen, autonome Unternehmen, die ihre eigenen Ziele verfolgen. Die tatsächliche Revolution dürfte prosaischer beginnen. Ein Planungssystem schlägt eine andere Reihenfolge für Fertigungsaufträge vor. Ein Vertriebsmodell setzt einen Preis, den der zuständige Manager für zu niedrig hält. Eine Software empfiehlt, einen langjährigen Lieferanten zugunsten eines bislang unauffälligen Wettbewerbers zurückzustufen. Niemand im Raum empfindet diese Entscheidungen als genial. Sie wirken nur irritierend. Einige Monate später zeigt die Ergebnisrechnung, dass sie richtig waren.

Solange solche Treffer vereinzelt bleiben, gelten sie als technische Kuriosität. Wiederholen sie sich, wird daraus eine Erwartung. Der entscheidende Übergang findet nicht im Modell, sondern in der Organisation statt: Ein Vorschlag wird vom Gegenstand menschlicher Prüfung zum Referenzpunkt, an dem sich das menschliche Urteil messen lassen muss. Aus der Frage „Ist die

KI gut genug, um uns zu beraten?“ wird die Frage „Können wir es uns leisten, ihren Rat zu ignorieren?“

Das Stockfish-Level bezeichnet genau diese Schwelle. Die Maschine ist nicht bloß schneller oder billiger. Sie trifft innerhalb eines definierten Feldes dauerhaft bessere Entscheidungen als die besten verfügbaren Menschen, und ihr Vorsprung ist groß genug, um Prozesse, Verantwortlichkeiten und Investitionslogiken zu verändern. Von diesem Moment an ist Delegation keine Modeerscheinung mehr. Sie wird zur ökonomischen Konsequenz.

Die Zahlen sind noch nicht Stockfish – aber sie zeigen die Richtung

Die bisherige empirische Forschung handelt überwiegend von Assistenzsystemen, nicht von übermenschlicher Entscheidungsautorität. Gerade deshalb ist sie aufschlussreich. In einem Feldversuch mit mehr als 5.000 Beschäftigten im Kundendienst stieg die Produktivität durch KI-Unterstützung im Durchschnitt um rund 14 Prozent; bei weniger erfahrenen und leistungsschwächeren Mitarbeitern betrug der Zuwachs etwa 34 Prozent. In kontrollierten Schreibaufgaben sank die benötigte Zeit um 40 Prozent, während die Qualität der Ergebnisse um 18 Prozent stieg. Und in einem Experiment mit Beratern wurden Aufgaben innerhalb der Leistungsgrenze des Systems mehr als 25 Prozent schneller und mit über 40 Prozent höher bewerteter Qualität bearbeitet.

Diese Ergebnisse sind kein Beweis dafür, dass heutige Systeme bereits das Stockfish-Level der Arbeitswelt erreicht haben. Sie zeigen vielmehr, wie schnell sich die ökonomische Beweislast verschiebt, sobald Leistungsunterschiede messbar werden. Ein Produktivitätsvorsprung von zwei oder drei Prozent kann in margenschwachen Branchen bereits strategisch relevant sein. Ein zweistelliger Unterschied verändert Personalbedarf, Durchlaufzeiten, Serviceversprechen und Preispositionen. Wer ihn dauerhaft nicht nutzt, entscheidet sich nicht für Menschlichkeit, sondern zunächst für höhere Kosten.

Bemerkenswert ist zudem, dass die Gewinne häufig bei weniger erfahrenen Beschäftigten besonders groß ausfallen. Die KI verdichtet das in der Organisation verteilte Wissen und macht es unmittelbar verfügbar. Das verkürzt Lernkurven und nivelliert Leistungsunterschiede. Für das Unternehmen ist das attraktiv; für die klassische Ordnung von Erfahrung und Karriere ist es Sprengstoff. Wenn ein Berufseinsteiger mit maschineller Unterstützung Ergebnisse erreicht, für die früher Jahre der Praxis nötig waren, verliert Erfahrung einen Teil ihres ökonomischen Monopols.

Vom Effizienzvorteil zur Wettbewerbsnorm

Neue Technologien verbreiten sich selten, weil alle Akteure von ihrem Sinn überzeugt sind. Sie verbreiten sich, weil einzelne Unternehmen mit ihnen Vorteile erzielen und damit die Bedingungen für alle anderen verändern. Das Fließband setzte sich nicht durch, weil es ästhetisch gefiel. Die Tabellenkalkulation wurde nicht eingeführt, weil Vorstände eine Leidenschaft für Zellen und Formeln entwickelten. Jede erfolgreiche Anwendung senkte Kosten, beschleunigte Entscheidungen oder erhöhte Qualität – und machte den bisherigen Standard relativ schlechter.

Beim Stockfish-Level wirkt derselbe Mechanismus mit größerer Geschwindigkeit. Ein Unternehmen, das Nachfrage, Wartung, Energieeinkauf und Lieferkettenrisiken besser prognostiziert, verfügt nicht nur über ein zusätzliches Werkzeug. Es kann niedrigere Bestände halten, Kapazitäten früher anpassen, knappe Ressourcen gezielter einsetzen und Angebote

schneller kalkulieren. Diese Vorteile verstärken einander. Bessere Entscheidungen führen zu besseren Daten, bessere Daten zu besseren Modellen, bessere Modelle zu noch schnelleren Entscheidungen. Der Vorsprung wird kumulativ.

Für Wettbewerber entsteht ein Delegationsdruck. Der Vorstand muss bald nicht mehr rechtfertigen, warum er eine maschinelle Empfehlung übernommen hat. Er muss erklären, warum er auf ein nachweislich überlegenes System verzichtete und dadurch einen Auftrag verlor, eine Störung zu spät erkannte oder Kapital fehlallokierte. In Haftungsdebatten wird dieser Punkt besonders sichtbar werden: Kann menschliches Urteil noch als Sorgfalt gelten, wenn vergleichbare Unternehmen systematisch bessere maschinelle Verfahren einsetzen?

Vier Quellen des ökonomischen Vorsprungs

Der wirtschaftliche Wert überlegener KI entsteht nicht nur durch geringere Personalkosten. Diese Verengung unterschätzt die eigentliche Dynamik. Erstens sinken Fehlerkosten. In komplexen Prozessen sind wenige vermiedene Fehlentscheidungen häufig wertvoller als tausend eingesparte Arbeitsstunden. Eine früh erkannte Qualitätsabweichung, ein vermiedener Lieferausfall oder eine richtig dimensionierte Wartung kann den Jahreswert eines Systems allein rechtfertigen.

Zweitens steigt die Entscheidungsgeschwindigkeit. Klassische Organisationen sammeln Daten, bereiten Vorlagen vor, koordinieren Abteilungen und warten auf Gremien. Ein System kann Szenarien fortlaufend aktualisieren und Entscheidungen innerhalb definierter Grenzen sofort auslösen. Der Wettbewerbsvorteil liegt dann nicht darin, dieselbe Antwort billiger zu produzieren, sondern Tage oder Wochen früher zu handeln.

Drittens sinken Koordinationskosten. Viele Managemententscheidungen sind weniger intellektuell schwierig als organisatorisch zäh. Vertrieb, Produktion, Einkauf und Finanzen optimieren unterschiedliche Ziele. Ein System kann ihre Wechselwirkungen simultan modellieren und Zielkonflikte sichtbar machen. Es ersetzt damit nicht nur einzelne Tätigkeiten, sondern einen Teil der Abstimmungsarbeit, aus der große Organisationen bestehen.

Viertens verbessert sich die Nutzung knapper Ressourcen. Kapital, Energie, Maschinenzeit und Fachkräfte können feiner zugeordnet werden. Kleine Prognosevorteile vervielfachen sich entlang der Wertschöpfungskette. Das erklärt, weshalb das Stockfish-Level nicht zuerst dort erfolgreich wird, wo spektakuläre Roboter auftreten, sondern dort, wo viele mittelgroße Entscheidungen täglich ineinandergreifen.

Die gezackte Grenze

Die ökonomische Versuchung zur Delegation trifft allerdings auf eine unangenehme Eigenschaft heutiger KI: Ihre Leistungsgrenze verläuft nicht sauber. Ein System kann eine komplexe Marktanalyse überzeugend erstellen und an einer banal formulierten Nebenbedingung scheitern. Es kann in hundert Fällen überlegen sein und im hundertersten mit großer Sicherheit Unsinn produzieren. Die Forschung spricht von einer „jagged technological frontier“, einer gezackten technologischen Grenze.

Für Unternehmen ist das gefährlicher als ein gleichmäßig mittelmäßiges System. Gleichmäßige Schwäche erzeugt Vorsicht; punktuelle Brillanz erzeugt Vertrauen. Je häufiger die KI richtigliegt, desto stärker wird die menschliche Neigung, auch jene Antworten zu akzeptieren, die außerhalb ihrer tatsächlichen Leistungszone liegen. Das Stockfish-Level kann deshalb nicht anhand einiger

eindrucksvoller Demonstrationen festgestellt werden. Es verlangt dauerhaft gemessene Überlegenheit in einer klar abgegrenzten Entscheidungsklasse.

Die wichtigste Investition ist folglich nicht das Modell allein, sondern die Vermessung seiner Grenze. Welche Aufgaben beherrscht es? Unter welchen Datenbedingungen? Wie verändert sich die Trefferquote bei neuartigen Situationen? Welche Fehler sind erkennbar, welche bleiben plausibel getarnt? Unternehmen, die diese Fragen nicht beantworten, besitzen keine überlegene Entscheidungsmaschine. Sie besitzen eine eindrucksvolle Blackbox mit unbekanntem Einsatzgebiet.

Die Kosten des menschlichen Vetos

Im traditionellen Management gilt das menschliche Veto als Ausdruck von Verantwortung. Doch ein Veto ist nicht kostenlos. Wenn Manager regelmäßig von einem überlegenen System abweichen, entstehen Opportunitätskosten. Sie bleiben meist unsichtbar, weil man nur den Schaden einer falschen KI-Empfehlung sieht, nicht aber den entgangenen Nutzen einer zu Unrecht verworfenen Empfehlung.

Deshalb brauchen Unternehmen eine neue Buchführung des Widerspruchs. Jede Abweichung sollte als überprüfbare Hypothese dokumentiert werden: Welche Information berücksichtigt das System angeblich nicht? Welcher Zielkonflikt rechtfertigt die Abweichung? Bis wann lässt sich das Ergebnis beurteilen? Erst dadurch wird das menschliche Veto von einer hierarchischen Geste zu einer lernfähigen Institution.

Das bedeutet nicht, Menschen durch Kennzahlen zu disziplinieren. Es schützt vielmehr den begründeten Widerspruch vor zwei schlechten Extremen: blindem Automatisierungsgelassenheit und managerialem Eigensinn. Wo Abweichungen weder begründet noch ausgewertet werden, gewinnt am Ende nicht die bessere Erkenntnis, sondern die mächtigere Position.

Nicht jede Entscheidung ist ein Schachzug

Die ökonomische Logik der Delegation gilt am stärksten dort, wo Ziele eindeutig, Ergebnisse zeitnah messbar und Entscheidungen reversibel sind. Bestandsoptimierung, Routenplanung, Maschinenbelegung oder Teile der Preissteuerung ähneln dem Schach vergleichsweise stark. Es gibt klare Regeln, wiederkehrende Situationen und genügend Daten, um Erfolg und Irrtum zu unterscheiden.

Standortschließungen, Personalentscheidungen oder strategische Richtungswechsel gehören in eine andere Klasse. Ihre Folgen entfalten sich über Jahre, Gegenwelten lassen sich nicht beobachten, und das Ergebnis hängt davon ab, was die Organisation überhaupt als Erfolg definiert. Eine KI kann hier bessere Szenarien liefern. Doch ihre statistische Überlegenheit beantwortet nicht, welche Interessen legitim sind und welche Lasten wem zugemutet werden dürfen.

Die wirtschaftliche Vernunft verlangt daher keine pauschale Automatisierung. Sie verlangt ein Portfolio: weitgehende Autonomie bei häufigen, reversiblen Entscheidungen; kontrollierte Delegation bei mittlerem Schadenspotenzial; Beratung ohne Letztentscheidung bei irreversiblen oder normativen Fragen. Gerade ein überlegenes System braucht eine präzise Grenze seiner Autorität.

Die Stockfish-Strategie

Eine tragfähige Strategie beginnt mit fünf nüchternen Fragen. Erstens: In welchen Entscheidungsklassen ist der Leistungsvorsprung wirklich gemessen? Zweitens: Welche wirtschaftlichen Effekte entstehen aus Qualität, Geschwindigkeit, Koordination und Ressourceneinsatz? Drittens: Wie teuer ist eine falsche maschinelle Entscheidung – und wie teuer ist das menschliche Ignorieren einer richtigen? Viertens: Wo liegt die gezackte Grenze des Systems? Fünftens: Welche Entscheidungen bleiben aus rechtlichen oder normativen Gründen menschlich, selbst wenn die Maschine statistisch überlegen erscheint?

Aus diesen Fragen folgt eine andere Investitionslogik. Unternehmen sollten nicht primär Modelle beschaffen, sondern Entscheidungsfähigkeiten aufbauen: Datenqualität, Vergleichstests, Protokolle für Abweichungen, unabhängige Kontrollen und ein klares Mandat für jede Entscheidungsklasse. Der Business Case einer KI besteht dann nicht in der Zahl erzeugter Texte, sondern in der nachweisbaren Verbesserung konkreter Entscheidungen.

Das Stockfish-Level beginnt wirtschaftlich dort, wo der Verzicht auf maschinelle Überlegenheit zum Wettbewerbsnachteil wird. Doch wer daraus bloß ein Automatisierungsprogramm macht, übersieht die eigentliche Zäsur. Die Firma kauft nicht nur Produktivität. Sie überträgt einem technischen System einen Teil ihrer Urteilsmacht. Das kann vernünftig sein. Es darf nur nicht unbemerkt geschehen.

Schluss: Der Preis des Nichtgebrauchs

Die kommende Auseinandersetzung wird nicht zwischen Unternehmen verlaufen, die KI lieben, und solchen, die sie ablehnen. Sie wird zwischen Organisationen verlaufen, die maschinelle Überlegenheit präzise messen und institutionell beherrschen, und solchen, die entweder euphorisch delegieren oder aus Gewohnheit widerstehen.

Auf dem Stockfish-Level wird das menschliche Urteil nicht wertlos. Es wird kostbarer, weil es seltener dort eingesetzt werden darf, wo die Maschine nachweislich besser ist, und gezielter dort gebraucht wird, wo Ziele, Grenzen und Legitimität zu bestimmen sind. Die neue Managementfrage lautet daher nicht: Mensch oder Maschine? Sie lautet: Wo ist der Preis des menschlichen Irrtums höher als der Preis der maschinellen Abhängigkeit – und wer darf diese Rechnung aufstellen?

Sobald diese Frage in Vorstandssitzungen ernsthaft gestellt wird, hat das Stockfish-Level die Arbeitswelt erreicht. Nicht weil die KI jedes Problem gelöst hätte. Sondern weil ihr Nichtgebrauch zu einer Entscheidung geworden ist, die sich rechtfertigen muss.

BEITRAG 2 | EPISTEMISCHE RESILIENZ

Die Organisation, die verlernt zu denken

Warum überlegene KI kurzfristig Leistung erzeugt und langfristig die menschliche Prüffähigkeit gefährden kann

Wenn eine Maschine regelmäßig besser entscheidet, ist es rational, ihr mehr zu überlassen. Genau darin liegt das Problem. Jede erfolgreiche Delegation verringert die Übung derjenigen, die das System im Ausnahmefall kontrollieren sollen. Die Organisation wird leistungsfähiger – und zugleich abhängiger von einer Intelligenz, die sie nicht mehr aus eigener Kraft ersetzen kann.

Die Abhängigkeitsspirale

Kurzfristiger Leistungsgewinn kann langfristig die menschliche Rückfallebene schwächen



Gegenmittel: epistemische Resilienz

Shadow Mode | regelmäßige manuelle Übungen | unabhängige Gegenmodelle | dokumentierte Overrides

INFPRO | Stockfish-Agenda | Konzeptionelles Modell

Abbildung 2: Konzeptionelles INFPRO-Modell.

Der Erfolg trägt das Risiko in sich

Schlechte Systeme erzeugen Widerstand. Gute Systeme erzeugen Gewohnheit. Sehr gute Systeme erzeugen Abhängigkeit. Diese einfache Folge erklärt, weshalb die größte organisatorische Gefahr nicht unbedingt von einer unzuverlässigen KI ausgeht. Ein unzuverlässiges System wird geprüft, diskutiert und notfalls abgeschaltet. Ein überlegenes System wird in den Alltag eingebaut. Es übernimmt immer mehr Entscheidungen, weil jede einzelne Übertragung vernünftig erscheint.

Damit entsteht eine Abhängigkeitsspirale. Bessere KI-Entscheidungen führen zu mehr Delegation. Mehr Delegation bedeutet weniger menschliche Übung. Weniger Übung senkt die Fähigkeit, ungewöhnliche Situationen zu beurteilen und Fehler zu erkennen. Sinkende

Prüffähigkeit macht die Organisation wiederum abhängiger vom System. Was auf jeder Stufe als Effizienzgewinn erscheint, kann im Ganzen zu einer strategischen Verwundbarkeit werden.

Diese Verwundbarkeit ist kein klassisches IT-Risiko. Server lassen sich spiegeln, Daten sichern, Anwendungen austauschen. Verlorene Urteilskompetenz lässt sich nicht kurzfristig aus einem Backup wiederherstellen. Sie verschwindet langsam, bleibt lange unsichtbar und fällt meist erst dann auf, wenn die Organisation sie dringend braucht.

Der Unterschied zwischen Wissen und Können

Unternehmen behandeln Wissen gern wie einen Bestand. Es wird dokumentiert, in Datenbanken gespeichert, in Prozessbeschreibungen überführt und durch Trainings verteilt. Doch ein erheblicher Teil professioneller Kompetenz ist kein abrufbarer Text. Er besteht aus Mustererkennung, situativem Urteil, eingeübter Aufmerksamkeit und dem Gespür dafür, wann eine scheinbar gewöhnliche Abweichung gefährlich wird.

Ein erfahrener Instandhalter hört nicht nur ein Geräusch. Er verbindet Klang, Lastzustand, Wartungshistorie und eine kaum bemerkte Veränderung im Verhalten der Anlage. Eine erfahrene Einkäuferin sieht nicht nur einen Preis. Sie erkennt, wann ein Lieferant ungewöhnlich bereitwillig zustimmt und welche Information ihm fehlen könnte. Solche Fähigkeiten entstehen durch wiederholte Entscheidungen mit Rückmeldung. Wer nur Ergebnisse bestätigt, ohne den Weg selbst zu gehen, sammelt weniger Erfahrung.

KI kann dieses implizite Wissen teilweise aufnehmen und skalieren. Genau das macht sie wertvoll. Doch wenn die Organisation das ursprüngliche menschliche Können nicht mehr reproduziert, verändert sich der Charakter des Wissens. Es ist vorhanden, aber nicht mehr verkörpert. Es arbeitet, aber niemand besitzt es. Das Unternehmen verfügt über Intelligenz, ohne sie selbst zu beherrschen.

Die verschwundene Lehrzeit

Besonders folgenreich ist dieser Effekt für den Nachwuchs. Experten entstehen nicht dadurch, dass sie fertige Antworten lesen. Sie entstehen durch Fälle, Irrtümer, Korrekturen und die allmähliche Verdichtung von Erfahrung. Wenn KI die schwierigen Zwischenschritte übernimmt, verbessert sie zwar die aktuelle Leistung der Anfänger, kann aber zugleich den Weg blockieren, auf dem aus Anfängern Experten werden.

Der Produktivitätsgewinn bei weniger erfahrenen Beschäftigten ist deshalb doppeldeutig. Er zeigt, wie wirksam KI Wissen verbreiten kann. Er wirft aber auch die Frage auf, ob der Mitarbeiter tatsächlich lernt oder nur vorübergehend auf die Kompetenz des Systems zugreift. Wer mit Navigation jedes Ziel findet, besitzt noch keine Ortskenntnis. Wer mit einer Diagnose-KI gute Treffer erzielt, hat noch nicht notwendigerweise gelernt, den atypischen Fall zu erkennen, in dem die Maschine scheitert.

Eine Organisation, die nur die unmittelbare Ergebnisqualität misst, übersieht diesen Unterschied. Sie kann heute mehr leisten und morgen weniger Menschen besitzen, die verstehen, warum. Damit wird der Talentprozess ausgehöhlt. Die Experten von gestern trainierten die Systeme; die Berufseinsteiger von heute bedienen sie; die Experten von morgen entstehen möglicherweise nicht mehr in ausreichender Zahl.

Kritisches Denken wird zur Resttätigkeit

Eine Untersuchung mit 319 Wissensarbeitern fand, dass höheres Vertrauen in generative KI mit weniger berichteter kritischer Prüfung verbunden war, während größeres Vertrauen in die eigene Fachkompetenz mit mehr Prüfung einherging. Die Studie beweist keinen zwangsläufigen geistigen Verfall. Sie zeigt jedoch einen plausiblen Mechanismus: Menschen verlagern ihre kognitive Anstrengung von der Erstellung auf die Kontrolle – und kontrollieren umso weniger, je verlässlicher die Maschine erscheint.

Das klingt zunächst effizient. Niemand prüft jede Formel einer Tabellenkalkulation von Hand. Doch die Kontrolle einer hochkomplexen Empfehlung verlangt häufig genau jene Sachkompetenz, die durch die Delegation seltener geübt wird. Man soll einen Fehler erkennen, ohne die Aufgabe selbst noch regelmäßig zu lösen. Das ist das Automationsparadox: Je zuverlässiger das System, desto seltener greift der Mensch ein; je seltener er eingreift, desto schlechter ist er auf den seltenen Moment vorbereitet, in dem sein Eingreifen zählt.

Aus kritischem Denken wird eine Resttätigkeit für Ausnahmefälle. Resttätigkeiten besitzen jedoch in Organisationen einen schlechten Stand. Sie werden nicht ausgelastet, erzeugen keine sichtbare Produktivität und lassen sich schwer in Budgets rechtfertigen. Gerade deshalb verschwinden sie – bis die Ausnahme eintritt.

Stockfish spielt in einer geschlossenen Welt

Im Schach ist die Abhängigkeit vergleichsweise harmlos. Das Brett hat 64 Felder, die Regeln bleiben stabil, die Figuren ändern ihre Eigenschaften nicht und niemand manipuliert während der Partie heimlich die Bedeutung des Sieges. Die Wirtschaft ist das Gegenteil einer solchen Welt. Kunden reagieren, Wettbewerber täuschen, Regierungen ändern Regeln, Lieferanten verschweigen Probleme und historische Daten verlieren nach einem Strukturbruch ihre Aussagekraft.

Ein System kann in der bekannten Welt überlegen sein und in einer veränderten Welt spektakulär scheitern. Es erkennt Muster, die bisher zuverlässig waren, und hält an ihnen fest, obwohl sich die Bedingungen verschoben haben. Der Mensch ist ebenfalls fehleranfällig, besitzt aber im besten Fall eine Fähigkeit, die sich schwer messen lässt: Er kann bemerken, dass die Lage nicht mehr zur Kategorie passt.

Das eigentliche Stockfish-Risiko liegt deshalb nicht im normalen Zug, sondern im unbemerkten Wechsel des Spiels. Eine KI optimiert weiterhin Schach, während die Organisation längst in einer Partie mit anderen Regeln steht. Wenn menschliche Fachleute ihre eigene Urteilspraxis verloren haben, fehlt gerade in diesem Moment die Instanz, die den Regimewechsel erkennt.

Die Luftfahrt als Warnung und Vorbild

Die Luftfahrt kennt das Problem seit Jahrzehnten. Automatisierung erhöht Sicherheit und entlastet Besatzungen, kann aber bei falscher Nutzung zu Verwirrung, Übervertrauen und schwindender manueller Routine führen. Leitlinien der amerikanischen Luftfahrtaufsicht betonen deshalb, dass manuelle Flugkompetenz die Grundlage weiterer fliegerischer Fähigkeiten bleibt und regelmäßig geübt werden muss. Moderne Systeme werden nicht aus Nostalgie zeitweise ausgeschaltet, sondern um die Rückfallebene real zu erhalten.

Für Unternehmen lässt sich daraus eine unbequeme Lektion ableiten. Eine Rückfallebene, die nur in einem Handbuch existiert, ist keine Rückfallebene. Menschen müssen Entscheidungen

ohne Systemunterstützung trainieren, Abweichungen erkennen und die Übernahme praktisch erproben. Das kostet Zeit und senkt kurzfristig die Effizienz. Doch auch Feuerübungen produzieren keinen Umsatz. Ihr Wert zeigt sich darin, dass die Organisation in einem seltenen kritischen Moment nicht erstmals über ihr Verhalten nachdenken muss.

Der Unterschied zur Luftfahrt ist allerdings erheblich: Für viele Managemententscheidungen gibt es weder standardisierte Simulatoren noch klar definierte Notfälle. Unternehmen müssen ihre Übungen selbst entwickeln. Sie brauchen Szenarien für Modell-Ausfälle, manipulierte Daten, widersprüchliche Empfehlungen, extreme Marktbrüche und den Verlust externer Anbieter. Epistemische Resilienz beginnt dort, wo diese Fälle nicht nur technisch, sondern als Entscheidungsproblem geprobt werden.

Was epistemische Resilienz bedeutet

Resilienz wird meist als Fähigkeit verstanden, nach einer Störung weiterzuarbeiten. Epistemische Resilienz geht tiefer. Sie bezeichnet die Fähigkeit einer Organisation, auch unter Bedingungen maschineller Überlegenheit eigene Urteils-, Prüf- und Handlungsfähigkeit zu bewahren. Nicht um jede Entscheidung besser als die KI zu treffen, sondern um ihre Grenzen zu erkennen, Ziele zu korrigieren und bei Ausfall oder Manipulation handlungsfähig zu bleiben.

Dazu gehören vier Reserven. Die erste ist eine Kompetenzreserve: Es müssen Menschen existieren, die Kernentscheidungen selbstständig treffen können. Die zweite ist eine Modellreserve: Kritische Systeme brauchen unabhängige Vergleichsverfahren, nicht bloß dieselbe Logik bei einem zweiten Anbieter. Die dritte ist eine Datenreserve: Originaldaten, Herkunft, Veränderungen und mögliche Manipulationen müssen nachvollziehbar bleiben. Die vierte ist eine institutionelle Reserve: Es braucht Rollen und Gremien, die begründeten Widerspruch nicht als Störung, sondern als Schutzfunktion behandeln.

Diese Reserven sind absichtlich redundant. Sie widersprechen dem klassischen Effizienzideal, nach dem jede Doppelstruktur als Verschwendung gilt. Doch vollständige Effizienz und echte Resilienz sind Gegensätze. Wer alle Reserven abbaut, optimiert das Unternehmen für eine Welt, in der nichts Unerwartetes geschieht.

Fünf Praktiken gegen das Verlernen

Erstens braucht es einen Shadow Mode. In ausgewählten Zeitfenstern lösen Menschen kritische Aufgaben unabhängig von der KI; die Ergebnisse werden erst anschließend verglichen. So lässt sich feststellen, ob Fachkompetenz erhalten bleibt und ob das System neue blinde Flecken entwickelt.

Zweitens sollten begründete Overrides systematisch ausgewertet werden. Nicht die Zahl der Widersprüche ist entscheidend, sondern ihre Qualität. Welche Gegenhypothesen waren richtig? Wo lag der Mensch daneben? Welche Signale hat das Modell übersehen? Aus dem Widerspruch entsteht nur dann Wissen, wenn die Organisation ihn erinnert.

Drittens müssen Lernwege neu gestaltet werden. Nachwuchskräfte dürfen nicht ausschließlich fertige KI-Lösungen freigeben. Sie müssen Fälle zunächst selbst strukturieren, Entscheidungen prognostizieren und erst danach die maschinelle Empfehlung sehen. Sonst lernen sie die Antwort, aber nicht das Urteil.

Viertens braucht es unabhängige Red Teams, die nicht nur Cyberangriffe simulieren, sondern die epistemische Stabilität prüfen: falsche Zielgrößen, verschobene Datenverteilungen, überzeugende Scheinerklärungen und koordinierte Modellfehler. Fünftens sollten Unternehmen

Ausfälle regelmäßig proben. Kann ein Werk 48 Stunden ohne Optimierungssystem produzieren? Kann die Kreditentscheidung bei einem Anbieter-Ausfall fortgeführt werden? Kann der Vorstand widersprüchliche Modelle einordnen, ohne sich die Entscheidung von einem dritten Modell abnehmen zu lassen?

Die neue Kennzahl: Wiederanlauffähigkeit des Urteils

Klassische Business Cases messen, was ein System zusätzlich leistet. Eine Stockfish-Governance muss auch messen, was ohne das System noch möglich ist. Dafür braucht es Kennzahlen: Anteil kritischer Prozesse mit trainierter manueller Rückfallebene; Zeit bis zur Wiederaufnahme des Betriebs ohne Modell; Quote der Entscheidungen, die Fachleute im Shadow Mode plausibel reproduzieren können; Zahl unabhängiger Daten- und Modellquellen; Qualität und Erfolgsrate begründeter Overrides.

Diese Kennzahlen werden nie so elegant sein wie Produktivität oder Kosten. Sie bewerten eine Fähigkeit, die im Normalbetrieb gerade nicht gebraucht wird. Doch darin besteht ihr Sinn. Epistemische Resilienz ist eine Versicherung gegen das Verschwinden der eigenen Urteilsmacht.

Das Ziel kann nicht sein, den Menschen künstlich auf dem Leistungsniveau der Maschine zu halten. Das wäre teuer und in vielen Bereichen aussichtslos. Er muss jedoch genug vom Problemraum verstehen, um einen Kategorienfehler zu erkennen, die Ziele des Systems zu hinterfragen und eine kontrollierte Notoperation zu ermöglichen.

Schluss: Intelligenz ohne Souveränität

Eine Organisation kann immer klüger entscheiden und zugleich immer weniger souverän werden. Dieser Satz beschreibt das Paradox des Stockfish-Levels. Solange das System funktioniert, erscheint Abhängigkeit als Effizienz. Erst im Ausfall, im Angriff oder im Strukturbruch zeigt sich, ob die Firma über eigene Urteilstkraft verfügt oder nur Zugang zu fremder Intelligenz hatte.

Die Antwort darf nicht in romantischer Handarbeit bestehen. Wer bessere Systeme aus Angst vor Kompetenzverlust nicht nutzt, konserviert nicht Urteilstkraft, sondern häufig nur schlechte Prozesse. Die Aufgabe ist anspruchsvoller: maschinelle Überlegenheit konsequent nutzen und zugleich die menschlichen Fähigkeiten erhalten, die man gerade deshalb immer seltener benötigt.

Das Unternehmen der Zukunft braucht daher nicht nur Datenresilienz, Cyberresilienz und Lieferkettenresilienz. Es braucht epistemische Resilienz. Andernfalls wird es vielleicht ausgezeichnet entscheiden – bis zu jenem Tag, an dem es erstmals wieder selbst denken muss.

BEITRAG 3 | ENTSCHEIDUNGSVERFASSUNG

Wer darf der Maschine widersprechen?

Wie Verantwortung, Haftung und Entscheidungsrechte auf dem Stockfish-Level neu geordnet werden müssen

Der Manager bleibt verantwortlich, obwohl er nicht mehr das beste Urteil besitzt. Er soll eine Entscheidung freigeben, die er weder entwickelt noch vollständig überprüfen kann. Das ist kein Randproblem der IT-Governance, sondern eine Verfassungskrise des Unternehmens: Wissen, Macht und Verantwortung fallen auseinander.

Decision Authority Matrix

KI-Autonomie folgt nicht der technischen Möglichkeit, sondern Reversibilität und Schadenspotenzial



INFPRO | Stockfish-Agenda | Konzeptionelles Modell

Abbildung 3: Konzeptionelles INFPRO-Modell.

Der verantwortliche Nichtwissende

Die moderne Organisation beruht auf einer stillen Verbindung. Wer entscheidet, soll wissen; wer weiß, soll Einfluss besitzen; wer Einfluss besitzt, soll Verantwortung tragen. In der Wirklichkeit war diese Ordnung nie vollkommen. Vorstände entscheiden über Technologien, die sie nicht im Detail verstehen, Ärzte verlassen sich auf Labore, Ingenieure auf Simulationen. Doch menschliche Experten konnten ihre Gründe erläutern, Gegenargumente aufnehmen und innerhalb einer gemeinsamen Fachsprache überzeugen.

Auf dem Stockfish-Level bricht diese Verbindung systematisch auseinander. Die KI verfügt über das bessere Urteil, besitzt aber weder Amt noch Haftung noch moralische Verantwortlichkeit. Der Manager besitzt Amt und Verantwortung, aber möglicherweise nicht mehr die Fähigkeit, die Entscheidung unabhängig zu prüfen. Er unterschreibt nicht sein eigenes Urteil. Er bestätigt ein Verfahren, dessen Überlegenheit statistisch belegt, dessen Einzelfalllogik jedoch nur begrenzt verständlich ist.

Damit entsteht die Figur des verantwortlichen Nichtwissenden. Sie ist nicht Ausdruck persönlicher Inkompetenz, sondern Folge einer strukturellen Asymmetrie. Selbst der beste menschliche Experte kann einem überlegenen System unterlegen sein. Die Governance muss deshalb nicht so tun, als ließe sich die alte Wissensordnung durch eine Unterschrift retten.

Human oversight ist noch keine menschliche Kontrolle

Der europäische AI Act verlangt für Hochrisiko-Systeme menschliche Aufsicht und nennt unter anderem die Fähigkeit, Systemgrenzen zu verstehen, Ausgaben zu interpretieren, Automatisierungsverzerrungen zu vermeiden, Eingriffe vorzunehmen und den Betrieb zu stoppen. Das ist ein wichtiger Rahmen. Doch zwischen formaler Aufsicht und wirksamer Kontrolle liegt eine große Lücke.

Ein Mensch im Prozess ist noch keine Kontrolle. Wenn die zuständige Person zu wenig Zeit, zu wenig Fachwissen oder keine realistische Möglichkeit zum Widerspruch besitzt, wird aus dem Human-in-the-loop ein dekorativer Freigabepunkt. Er bestätigt, was das System vorgibt, und verleiht der Entscheidung den Anschein persönlicher Verantwortung. Die Organisation erhält damit Haftungsfolklore statt Governance.

Wirksame Aufsicht verlangt drei Fähigkeiten: verstehen, worüber entschieden wird; erkennen, wann die Situation außerhalb des vorgesehenen Einsatzbereichs liegt; und widersprechen können, ohne dass der Prozess oder die Karriere den Widerspruch bestraft. Fehlt eine dieser Bedingungen, ist der Mensch kein Kontrolleur, sondern eine menschliche Signatur unter einer maschinellen Entscheidung.

Entscheidungsrechte statt Modellfreigaben

Viele Unternehmen organisieren KI-Governance entlang von Modellen. Ein System wird klassifiziert, getestet, genehmigt und überwacht. Das ist notwendig, aber nicht ausreichend. Dasselbe Modell kann eine harmlose Textzusammenfassung erstellen oder eine folgenschwere Personalentscheidung vorbereiten. Das Risiko steckt nicht allein im Modell, sondern in der Kombination aus Entscheidung, Einsatzkontext und Autorität.

Die zentrale Einheit der Governance muss deshalb die Entscheidung sein. Für jede relevante Entscheidungsklasse ist festzulegen: Welches Ziel wird optimiert? Welche Daten dürfen verwendet werden? Wie reversibel ist das Ergebnis? Welches Schadenspotenzial besteht? Wie schnell muss gehandelt werden? Wer darf zustimmen, wer widersprechen, wer stoppen? Welche Beweise sind vor und nach der Entscheidung erforderlich?

Erst daraus entsteht eine Decision Authority Matrix. Sie ordnet nicht abstrakt Menschen und Maschinen, sondern konkrete Rechte: empfehlen, entscheiden unter Vorbehalt, autonom innerhalb von Grenzen handeln oder vom Einsatz ausgeschlossen bleiben. Diese Rechte müssen enger sein als die technische Fähigkeit des Systems. Eine Maschine darf nicht alles entscheiden, was sie entscheiden kann.

Reversibilität ist die unterschätzte Achse

Die Debatte über KI-Risiken konzentriert sich häufig auf Eintrittswahrscheinlichkeiten. Für Governance ist jedoch die Reversibilität mindestens ebenso wichtig. Eine falsche Lagerdisposition kann korrigiert werden. Eine falsch ausgelöste Sicherheitsabschaltung ist teuer, aber oft begrenzt. Eine Standortschließung, eine medizinische Intervention oder die öffentliche Beschuldigung eines Menschen lassen sich dagegen nicht vollständig zurücknehmen.

Je irreversibler eine Entscheidung und je höher ihr Schadenspotenzial, desto weniger darf statistische Überlegenheit allein ihre Legitimität begründen. Das ist kein Misstrauensvotum gegen KI, sondern eine allgemeine Regel vernünftiger Institutionen. Auch menschliche Experten dürfen in solchen Bereichen nicht grenzenlos entscheiden. Sie unterliegen Verfahren, Anhörungen, Dokumentationspflichten und mehreren Instanzen.

Für hochreversible Entscheidungen mit begrenztem Schaden kann weitgehende Autonomie sinnvoll sein. Bei mittlerer Reversibilität braucht es Echtzeitmonitoring, Stop-Regeln und ex-post Kontrolle. Bei folgenreichen, kaum reversiblen Entscheidungen sollte die KI Szenarien, Prognosen und Gegenargumente liefern, aber nicht die normative Letztentscheidung besitzen.

Das Human Override Protocol

Ein menschliches Veto muss möglich sein, darf aber nicht beliebig werden. Sonst entsteht eine merkwürdige Organisation: Offiziell gilt die KI als überlegen; praktisch ignoriert jede Führungskraft ihre Empfehlungen nach Geschmack. Umgekehrt darf die Erfolgsbilanz des Systems nicht zu einem Klima führen, in dem niemand mehr zu widersprechen wagt.

Ein belastbares Override-Protokoll verlangt zunächst einen benannten Grund. Zulässig sind beispielsweise neue Informationen außerhalb des Datenstands, ein erkennbarer Regimewechsel, ein Konflikt mit einer unverhandelbaren Grenze, eine unvertretbare Verteilung von Schäden oder ein begründeter Verdacht auf Manipulation. Das Veto sollte eine Gegenhypothese formulieren, eine zweite fachliche Prüfung auslösen und einen Zeitpunkt für die Erfolgskontrolle bestimmen.

Wichtig ist die Schutzfunktion. Wer begründet widerspricht, darf nicht schon deshalb sanktioniert werden, weil die Maschine in der Mehrzahl der Fälle recht behält. Sonst wird das Override zu einem roten Knopf hinter Glas, den niemand anzufassen wagt. Gleichzeitig müssen wiederholt unbegründete Abweichungen sichtbar werden. Widerspruch ist ein Recht – und eine überprüfbare Verantwortung.

Wenn die Erklärung nur eine Geschichte ist

Erklärbarkeit gilt als Lösung des Kontrollproblems. Systeme sollen Einflussfaktoren zeigen, Begründungen liefern und alternative Szenarien darstellen. Das verbessert die Aufsicht, kann aber eine gefährliche Illusion erzeugen. Eine verständliche Erklärung ist nicht notwendigerweise die tatsächliche Ursache einer Entscheidung. Sie kann eine vereinfachte Übersetzung oder eine nachträgliche Rationalisierung sein.

Je größer der Abstand zwischen maschineller und menschlicher Leistungsfähigkeit, desto stärker dieser Konflikt. Eine Erklärung, die vollständig auf menschliches Niveau reduziert wird, kann gerade jene komplexen Wechselwirkungen unterschlagen, aus denen der Leistungsvorsprung entsteht. Umgekehrt ist eine technisch vollständige Beschreibung wertlos, wenn kein Verantwortlicher sie in der verfügbaren Zeit verstehen kann.

Governance darf deshalb Erklärbarkeit nicht mit Verifikation verwechseln. Eine Entscheidung kann erklärbar und falsch, unerklärbar und richtig oder gut begründet, aber auf ein falsches Ziel optimiert sein. Benötigt werden mehrere Evidenzarten: Leistungsdaten, Robustheitstests, Gegenmodelle, Sensitivitätsanalysen, Herkunft der Daten und eine klare Darstellung dessen, was unbekannt bleibt.

Die Zielfunktion ist eine politische Entscheidung

Das gefährlichste System ist nicht unbedingt eines, das schlecht optimiert. Es kann eines sein, das ein unvollständiges Ziel hervorragend verfolgt. Wer Durchlaufzeit minimiert, kann Reserven zerstören. Wer Personalkosten senkt, kann Wissen verlieren. Wer Ausfallwahrscheinlichkeit reduziert, kann Innovation verhindern. Auf dem Stockfish-Level werden falsch gesetzte Ziele nicht milder, sondern konsequenter umgesetzt.

Zieldefinition ist deshalb keine technische Vorarbeit, die an Datenwissenschaftler delegiert werden kann. Sie ist eine Führungs- und häufig eine politische Aufgabe. Welche Größen sollen verbessert werden? Welche dürfen nur innerhalb von Korridoren schwanken? Welche Nebenbedingungen sind unverhandelbar? Wie werden Interessen gewichtet, die sich nicht in derselben Einheit ausdrücken lassen?

Der Vorstand wird damit weniger zum besten Spieler und stärker zum Verfassungsgeber. Er legt fest, welche Partien gespielt, welche Ziele verfolgt und welche Züge ausgeschlossen werden. Diese Rolle ist nicht geringer als die traditionelle Entscheidungsmacht. Sie ist anspruchsvoller, weil sie die Bedingungen für tausende maschinelle Einzelentscheidungen setzt.

Haftung muss am Prozess ansetzen

Die klassische Haftungsintuition sucht eine Person: Wer hat entschieden? Auf dem Stockfish-Level ist diese Frage zu grob. Eine maschinelle Entscheidung entsteht aus Zielvorgaben, Daten, Modellen, Systemintegration, Freigaben, Überwachung und organisatorischem Druck. Die Verantwortung verschwindet nicht, aber sie verteilt sich entlang dieser Kette.

Ein Manager kann nicht sinnvoll für jede interne Modelloperation haften. Er kann jedoch dafür verantwortlich sein, dass die Entscheidungsklasse korrekt eingeordnet, der Einsatzbereich definiert, die Leistungsfähigkeit belegt, das Override funktionsfähig und die Nachkontrolle organisiert ist. Verantwortung verlagert sich vom Besitz des besseren Einzelurteils zur Qualität des Entscheidungssystems.

Das NIST AI Risk Management Framework fasst Risikomanagement in die Funktionen Govern, Map, Measure und Manage. Für die Unternehmenspraxis ist diese Logik hilfreich: Zuständigkeiten festlegen; Kontext und Folgen der Entscheidung verstehen; Leistung und Risiken messen; Grenzen, Störungen und Verbesserungen aktiv managen. Entscheidend ist, dass daraus keine Checkliste ohne Macht folgt. Wer das System stoppen soll, braucht tatsächlich die Befugnis dazu.

Die Aufgaben von Vorstand und Aufsichtsrat

Der Vorstand muss zunächst ein Inventar wesentlicher KI-gestützter Entscheidungen schaffen. Nicht Anwendungen zählen, sondern Entscheidungen mit Wirkung: Preise, Kredite, Personal, Wartung, Investitionen, Sicherheit, Beschaffung, Produktion. Für jede Klasse sind Autonomiegrad, Schadenspotenzial, Reversibilität und Rückfallebene festzulegen.

Der Aufsichtsrat sollte sich nicht mit allgemeinen Berichten über „KI-Initiativen“ zufriedengeben. Er braucht eine Sicht auf die größten Entscheidungsabhängigkeiten: Wo kann kein Mensch mehr unabhängig prüfen? Wo existiert nur ein Modell oder Anbieter? Welche Ziele wurden vom Vorstand ausdrücklich legitimiert? Wie oft wurde widersprochen, mit welchem Ergebnis? Wann wurde der letzte Ausfall geprobt?

Daraus entsteht eine neue Form der Assurance. Interne Revision, Risikomanagement, Compliance und Fachbereiche müssen nicht das Modell im selben Sinne verstehen wie seine Entwickler. Sie müssen gemeinsam nachweisen können, dass die Autorität des Systems begrenzt, beobachtet und im Störfall zurückgenommen werden kann.

Ein 90-Tage-Programm

In den ersten 30 Tagen sollte das Unternehmen seine entscheidungsrelevanten KI-Einsätze erfassen und nach Reversibilität sowie Schadenspotenzial klassifizieren. Gleichzeitig ist festzustellen, wo menschliche Entscheidungen bereits nur noch formal bestätigt werden. Diese Schatten-Delegation ist häufig riskanter als eine offen geregelte Autonomie.

Bis Tag 60 werden für die wichtigsten Entscheidungsklassen Leistungsbaselines, zulässige Autonomiegrade und Override-Gründe definiert. Es geht nicht um perfekte Vollständigkeit, sondern um die kritischen zwanzig Prozent der Entscheidungen, die achtzig Prozent des Risikos tragen. Parallel werden unabhängige Kontrollen und Rückfallverfahren benannt.

Bis Tag 90 sollten Vorstand und Aufsichtsrat die Decision Authority Matrix verabschieden, einen ersten Ausfall- oder Regimewechseltest durchführen und ein Berichtssystem einrichten. Es zeigt nicht nur Modellfehler, sondern auch Kosten menschlicher Abweichungen, Kompetenzverluste und offene Zielkonflikte. Danach beginnt die eigentliche Arbeit: Governance muss mit dem Leistungszuwachs der Systeme Schritt halten.

Schluss: Die Verfassung vor dem Tempo

Auf dem Stockfish-Level kann Führung nicht mehr darin bestehen, bei jeder Einzelfrage das beste Urteil zu beanspruchen. Eine solche Pose wäre weder glaubwürdig noch wirtschaftlich vernünftig. Führung besteht darin, die Autorität überlegener Systeme zu begrenzen, ohne ihre Erkenntnisse aus gekränkter Selbstachtung zu verschwenden.

Dafür braucht das Unternehmen eine Entscheidungsverfassung. Sie regelt Ziele, Rechte, Grenzen, Beweise, Widerspruch und den Ausnahmezustand. Sie schützt Menschen nicht vor jeder maschinellen Entscheidung, sondern vor einer Macht, die sich allein aus Trefferquoten legitimiert. Und sie schützt die Organisation vor Managern, die ihre Hierarchie mit Urteilskraft verwechseln.

Die Frage lautet am Ende nicht, ob der Mensch der Maschine widersprechen darf. Sie lautet: Unter welchen Bedingungen ist sein Widerspruch eine verantwortliche Entscheidung – und unter welchen Bedingungen nur die teuerste Form der Eitelkeit? Wer diese Grenze nicht institutionell zieht, wird sie im Ernstfall improvisieren. Dann entscheidet nicht die klügere Instanz, sondern die lautere.

QUELLEN UND HINWEISE

Empirische Ankerpunkte

1. Brynjolfsson, Li & Raymond (2023/2025)

Generative AI at Work. Feldstudie mit über 5.000 Kundendienstmitarbeitern; Produktivitätszuwachs im Durchschnitt rund 14 Prozent, besonders stark bei weniger erfahrenen Beschäftigten.

<https://www.nber.org/papers/w31161>

2. Noy & Zhang (Science, 2023)

Experimental Evidence on the Productivity Effects of Generative Artificial Intelligence. In professionellen Schreibaufgaben sank die Bearbeitungszeit um 40 Prozent, die Qualität stieg um 18 Prozent.

<https://www.science.org/doi/10.1126/science.adh2586>

3. Dell'Acqua et al. (Harvard Business School, 2023)

Navigating the Jagged Technological Frontier. Feldexperiment mit Beratern; deutliche Gewinne innerhalb, aber Risiken außerhalb der Leistungsgrenze des Systems.

<https://www.hbs.edu/faculty/Pages/item.aspx?num=64700>

4. Lee et al. (CHI/Microsoft Research, 2025)

The Impact of Generative AI on Critical Thinking. Befragung von 319 Wissensarbeitern; höheres Vertrauen in GenAI war mit weniger berichteter kritischer Prüfung verbunden.

<https://www.microsoft.com/en-us/research/publication/the-impact-of-generative-ai-on-critical-thinking-self-reported-reductions-in-cognitive-effort-and-confidence-effects-from-a-survey-of-knowledge-workers/>

5. Europäische Union (2024)

Verordnung (EU) 2024/1689, insbesondere Artikel 14 zur menschlichen Aufsicht über Hochrisiko-KI-Systeme.

<https://eur-lex.europa.eu/eli/reg/2024/1689/oj>

6. NIST (2023/2024)

AI Risk Management Framework und GenAI Profile. Governance-Logik mit den Funktionen Govern, Map, Measure und Manage.

<https://www.nist.gov/itl/ai-risk-management-framework>

7. Federal Aviation Administration (2022)

Advisory Circular AC 120-123. Manuelle Flugkompetenz als Grundlage und notwendige Rückfallebene in automatisierten Cockpits.

https://www.faa.gov/documentLibrary/media/Advisory_Circular/AC_120-123.pdf

Redaktioneller Hinweis

Die Studien belegen Leistungs- und Verhaltenswirkungen heutiger Assistenzsysteme. Der Begriff „Stockfish-Level“ und die daraus abgeleiteten Modelle sind eine konzeptionelle Fortschreibung: Sie beschreiben einen möglichen Schwellenzustand überlegener, nur noch begrenzt menschlich verifizierbarer KI. Empirische Befunde und Zukunftsszenario werden deshalb bewusst getrennt.