



SITUATIONAL AWARENESS

The Decade Ahead

Leopold Aschenbrenner — Juni 2024

Zusammenfassung & Kernanalyse

Golden Circle Finance · Interne Analyse · Juni 2026

*AGI-Awareness: Nur wenige Hundert Menschen weltweit verstehen die Trendlinien.
AGI bis 2027 ist keine Science-Fiction -- es ist einfache Extrapolation.*

Executive Summary

Leopold Aschenbrenner, ehemaliger Safety-Researcher bei OpenAI, hat mit 'Situational Awareness: The Decade Ahead' (Juni 2024) eine der detailliertesten und am meisten zitierten Analysen der gegenwärtigen KI-Entwicklung vorgelegt. Das Dokument umfasst 165 Seiten und gliedert sich in fünf Hauptkapitel, die gemeinsam ein kohärentes Bild zeichnen: Die Welt steht am Vorabend einer technologischen Revolution, die alles bisher Dagewesene in den Schatten stellen wird.

Aschenbrenners Kernthese: Allgemeine Künstliche Intelligenz (AGI – Artificial General Intelligence) — KI-Systeme, die die kognitive Leistung von Hochschulabsolventen übertreffen und im Wesentlichen jede Wissensarbeit automatisieren können — ist bis 2027 wahrscheinlich. Danach folgt ein rapider Sprung zu Superintelligenz, möglicherweise innerhalb von nur einem weiteren Jahr. Diese Entwicklung wird geopolitische, militärische und wirtschaftliche Auswirkungen mit sich bringen, die historisch beispiellos sind.



Über den Autoren:

Leopold Aschenbrenner (geboren 2001 oder 2002) ist ein deutscher Forscher und Investor im Bereich der Künstlichen Intelligenz (KI). Er gilt in der Tech- und Finanzwelt als eines der einflussreichsten „Wunderkinder“ rund um die zukünftige Entwicklung und Monetarisierung von KI-Technologien.

Herkunft und Ausbildung

- **Frühes Leben:** Er wuchs in Deutschland als Sohn eines Ärzteehepaares auf und besuchte die John-F.-Kennedy-Schule in Berlin.
- **Studium:** Er beendete 2021 im Alter von nur 19 Jahren sein Studium an der Columbia University als Jahrgangsbester. Er hält Abschlüsse in Wirtschaftswissenschaften sowie Mathematik und Statistik.
- **Effektiver Altruismus:** Während des Studiums engagierte er sich stark in der Bewegung des Effektiven Altruismus und arbeitete zeitweise für den FTX Future Fund.

Karriere bei OpenAI und Entlassung

Aschenbrenner stieß 2023 zu OpenAI und wurde Mitglied des „Superalignment“-Teams unter der Leitung von Ilya Sutskever und Jan Leike. Dieses Team befasste sich mit der Kontrolle und Sicherheit von zukünftigen KI-Systemen, die klüger als Menschen sind.

Im April 2024 wurde er von OpenAI entlassen. Der offizielle Grund war eine angebliche Weitergabe von Informationen. Aschenbrenner selbst gab an, dass seine Entlassung politisch motiviert war, nachdem er ein internes Warn-Memo bezüglich unzureichender IT-Sicherheit und der Gefahr von Industriespionage (insbesondere durch China) an den Vorstand geschickt hatte.

Der Essay „Situational Awareness“

Kurz nach seinem Ausscheiden im Juni 2024 veröffentlichte er ein vielbeachtetes, 165-seitiges Papier namens „Situational Awareness: The Decade Ahead“. Dieser Essay verbreitete sich rasant im Silicon Valley und stellt folgende Thesen auf:

- Eine allgemeine künstliche Intelligenz (AGI), die menschliche Fähigkeiten übertrifft, ist bereits bis **2027** realistisch.
- Das KI-Wachstum verläuft sprunghaft und wird zu Billionen-Dollar-Rechenclustern führen.
- Der größte Flaschenhals der nahen Zukunft ist nicht die Software oder die Chips, sondern die **Energie- und Stromversorgung**.



I. Von GPT-4 zu AGI: Die OOM-Analyse

Das Tempo des Fortschritts

Aschenbrenners Ausgangspunkt ist eine scheinbar simple Beobachtung: Von GPT-2 (2019) zu GPT-4 (2023) hat die KI in nur vier Jahren einen Entwicklungssprung vollzogen, der sich im menschlichen Kontext als Weg vom Vorschulkind zum intelligenten Oberschüler beschreiben lässt. GPT-2 konnte kaum kohärente Sätze bilden. GPT-4 besteht anspruchsvolle Hochschulprüfungen, schreibt komplexen Code und denkt durch schwierige mathematische Probleme.

Die Frage lautet: Warum sollte sich dieser Trend nicht fortsetzen? Aschenbrenner argumentiert überzeugend, dass er es wird — und analysiert drei fundamentale Treiber dieser Entwicklung in sogenannten OOMs (Orders of Magnitude, Größenordnungen):

- **Rechenleistung (Compute):** Zwischen GPT-2 und GPT-4 wuchs die eingesetzte Rechenleistung um das 3.000- bis 10.000-fache. Bis Ende 2027 sind weitere 100- bis 1.000-fache Steigerungen realistisch — getrieben durch Investitionen in Milliarden- und Billionenhöhe.
- **Algorithmische Effizienz:** Parallel zu mehr Hardware wird die Software effizienter. Modelle, die gleiche Leistung mit 1.000-mal weniger Rechenleistung erbringen, sind bereits Realität. Der Trend zeigt ca. 0,5 OOM Verbesserung pro Jahr, ein Faktor von rund 10-facher Steigerung alle vier Jahre.
- **'Unhobbling' (Befreiung latenter Fähigkeiten):** Aktuelle Modelle sind künstlich eingeschränkt. Chain-of-Thought-Prompting, RLHF, Scaffolding und Werkzeugintegration haben bereits enorme Leistungssprünge ermöglicht. Noch sind die Modelle weit von ihrem Potenzial entfernt.

Der Datenwand-Einwand

Aschenbrenner adressiert das häufigste Gegenargument: das baldige Erschöpfen verfügbarer Trainingsdaten. Er hält dies für überwindbar — durch synthetische Daten, Self-Play und Reinforcement Learning (vergleichbar mit dem Weg von AlphaGo). Er räumt ein, dass hier der größte Unsicherheitsfaktor liegt, geht aber davon aus, dass die Labs diesen Flaschenhals überwinden werden.

Das Ergebnis: AGI by 2027

Die Schlussfolgerung: Bis Ende 2027 werden KI-Modelle existieren, die das Wissen und die Problemlösungsfähigkeit von KI-Forschern und Ingenieuren vollständig abdecken. Nicht als Chatbot, sondern als autonomer Agent -- ein sogenannter Drop-in Remote Worker, der in Unternehmen eingesetzt werden kann wie ein Mitarbeiter, der von überall auf der Welt arbeitet.

'Wir sind dabei, die OOMs so schnell zu durchrasen, dass es keines esoterischen Glaubens bedarf, um AGI bis 2027 ernst zu nehmen — nur Trendextrapolation auf geraden Linien.'



II. Von AGI zu Superintelligenz: Die Intelligenzexplosion

Warum KI-Fortschritt nicht bei menschlichem Niveau aufhört

Sobald AGI erreicht ist, werden wir nicht einfach ein sehr gutes KI-Werkzeug haben. Wir werden die Möglichkeit haben, KI-Forschung selbst zu automatisieren. Das ist der zentrale, oft unterschätzte Schritt: Mit Millionen von AGI-Instanzen, die 24/7 laufen und algorithmische Forschung betreiben, könnte sich ein Jahrzehnt menschlicher KI-Forschung in einem einzigen Jahr vollziehen.

Aschenbrenner schätzt, dass bis 2027 GPU-Flotten existieren werden, die den Betrieb von mindestens 100 Millionen 'Human-Equivalent-Researcher' ermöglichen — und das bei Geschwindigkeiten, die um ein Vielfaches über menschlichen Denkgeschwindigkeiten liegen. Das Ergebnis: 5+ OOM algorithmische Effizienzgewinne innerhalb eines Jahres, was einem weiteren GPT-2-zu-GPT-4-Sprung entspräche — diesmal auf Basis von AGI.

Mögliche Flaschenhälse

Der Autor diskutiert ehrlich die Gegenkräfte: begrenzte Rechenkapazität für Experimente, fehlende komplementäre menschliche Fähigkeiten, das Schwieriger-Werden von Innovationen. Er kommt dennoch zu dem Schluss, dass ein Zeitrahmen von einem bis wenigen Jahren für den Übergang zu echter Superintelligenz realistisch ist.

Die Macht der Superintelligenz

Superintelligenz wäre nicht nur quantitativ überlegen (mehr Wissen, mehr Geschwindigkeit), sondern qualitativ anders — vergleichbar mit dem qualitativen Unterschied zwischen einem Erstklässler und einem promovierten Wissenschaftler. Solche Systeme würden:

- wissenschaftlichen Fortschritt in Dekaden in Jahren vollziehen
- Robotik- und Industrieautomatisierung lösen
- wirtschaftliches Wachstum von 30% pro Jahr oder mehr ermöglichen
- entscheidende militärische Überlegenheit verschaffen — sogar gegen nukleare Abschreckung
- potenziell die Machtbalance ganzer Zivilisationen verschieben

'Die Intelligenzexplosion wird eine der volatilsten, intensivsten und gefährlichsten Perioden in der Geschichte der Menschheit sein. Und bis Ende des Jahrzehnts werden wir mitten darin sein.'



III. Die vier großen Herausforderungen

IIIa. Das Billionen-Dollar-Cluster

Die KI-Revolution ist auch eine Industrie- und Kapitalrevolte. Aschenbrenner zeichnet nach, wie sich die Investitionen in KI-Infrastruktur in astronomische Größenordnungen entwickeln:

- 2024: ~150 Milliarden Dollar Gesamtinvestition in KI
- 2026: ~500 Milliarden Dollar
- 2028: ~2 Billionen Dollar
- 2030: ~8 Billionen Dollar

Das Training des größten Modells benötigt bis 2030 voraussichtlich 100 GW Strom — mehr als 20% des gesamten US-amerikanischen Stromverbrauchs. Microsoft und OpenAI planen laut Medienberichten bereits einen 100-Milliarden-Dollar-Cluster für 2028. Die zentrale Herausforderung ist nicht das Kapital, sondern die Infrastruktur: Strom, Land, Genehmigungen.

Aschenbrenner plädiert dafür, diese Cluster aus geopolitischen Gründen in den USA oder demokratischen Verbündeten zu bauen — nicht in autoritären Staaten, die das Kapital anbieten, aber die strategische Kontrolle übernehmen würden.

IIIb. Sicherheit der KI-Labors

Dies ist möglicherweise das dringlichste und gleichzeitig am meisten vernachlässigte Thema des gesamten Dokuments. Aschenbrenner argumentiert schonungslos: Die führenden KI-Labors behandeln Sicherheit wie ein beliebiges Tech-Startup — obwohl sie die strategisch wichtigsten Technologiesecrets der USA entwickeln.

Die zwei kritischen Angriffsvektoren:

- Modellgewichte (Weights): Der fertige KI-Code ist eine Datei auf einem Server. Wird diese gestohlen, spart China Billionen an Investitionen und Jahre Rückstand. Wer die AGI-Weights hat, hat Superintelligenz.
- Algorithmische Secrets: Die Forschungsdurchbrüche, die auf dem Weg zu AGI entstehen (z.B. neue Methoden zur Überwindung der Datenwand), sind möglicherweise noch wertvoller. Sie können in einem kurzen Gespräch übermittelt werden — und sind derzeit kaum geschützt.

Aschenbrenner zitiert Marc Andreessen: Er gehe davon aus, dass die chinesische Regierung bereits vollständigen Zugang zu amerikanischer KI-Forschung habe. Die Konsequenz: Wenn die USA ihre KI-Geheimnisse nicht sofort und radikal sichern, verschaffen sie China de facto den Schlüssel zur Superintelligenz.



IIIc. Superalignment — KI-Kontrolle

Das technische Problem: Wie kontrollieren wir KI-Systeme, die klüger sind als wir? RLHF (Reinforcement Learning from Human Feedback) — die aktuelle Standardmethode — bricht zusammen, sobald KI menschliche Kompetenz überschreitet. Ein Mensch kann nicht beurteilen, ob eine Million Zeilen Code in einer neu erfundenen Programmiersprache Sicherheitslücken enthält.

Aschenbrenner skizziert einen Ausweg: schrittweise Kontrolle durch

- Scalable Oversight: KI hilft Menschen, andere KI zu überwachen
- Generalisierung: Korrekte Werte aus einfachen Fällen sollen auf komplexe generalisieren
- Interpretierbarkeit: Verständnis der internen Prozesse
- Adversariales Testen: Systematische Suche nach Fehlern

Er ist vorsichtig optimistisch, dass diese Probleme lösbar sind — aber alarmiert davon, wie wenig ernsthafte Arbeit derzeit daran geleistet wird. Die Intelligenzexplosion macht das Problem besonders kritisch: In unter einem Jahr könnte man von kontrollierbaren zu unkontrollierbaren Systemen übergehen.

IIId. Der geopolitische Kampf: Die freie Welt muss gewinnen

Superintelligenz verleiht denjenigen, die sie zuerst besitzen, einen entscheidenden militärischen Vorteil — vergleichbar mit dem Übergang von konventionellen Waffen zu Atomwaffen, nur noch dramatischer. Aschenbrenner analysiert die chinesische KI-Position nüchtern:

- China kann 7nm-Chips produzieren (vergleichbar mit Nvidia A100s)
- China baut Infrastruktur schneller als die USA (Strom, Rechenzentren)
- China betreibt massiv industrielle Spionage gegen US-KI-Labors
- China könnte AGI-Weights stehlen und eine eigene Intelligenzexplosion starten

Das autoritäre Risiko: Wer mit Superintelligenz regiert, kann Macht auf eine Art konsolidieren, die in der Geschichte beispiellos ist — mit KI-gesteuerter Massenüberwachung, autonomen Robotermilizen und absoluter Kontrolle über Information. Aschenbrenner nennt das 'Value Lock-in' und sieht es als existenzielle Bedrohung für Demokratie und Freiheit.

'Die freie Welt muss im Rennen um AGI siegen. Wir verdanken unseren Frieden und unsere Freiheit der wirtschaftlichen und militärischen Überlegenheit Amerikas. Wenn wir dieses Rennen verlieren, riskieren wir alles.'



IV. Das Projekt — staatliche Übernahme unvermeidlich

Aschenbrenners mutigste These: Die US-Regierung wird in den nächsten Jahren notgedrungen die Kontrolle über die führenden KI-Entwicklungen übernehmen. Nicht weil es eine gute Idee ist, sondern weil es unvermeidlich ist. Er nennt es schlicht 'The Project' — in Anlehnung an das Manhattan Project.

Der Weg dorthin

Aschenbrenner skizziert einen Zeitplan: Um 2026/27 werden KI-Demos existieren, die jeden Beobachter erschüttern — vielleicht die Unterstützung bei Biowaffen-Entwicklung, autonomes Hacking kritischer Infrastruktur oder ähnliches. Die Geheimdienstgemeinschaft wird bemerken, dass China massiv in amerikanische KI-Labors infiltriert hat. Die öffentliche Debatte wird kippen. Der Kongress wird Billionen bewilligen.

Das Projekt wird keine vollständige Verstaatlichung sein, aber eine enge Kooperation zwischen führenden Tech-Unternehmen, Verteidigungsministerium und Geheimdiensten — ähnlich dem Verhältnis zwischen Pentagon und Boeing oder Lockheed Martin. Die führenden KI-Labors werden sich 'freiwillig' zusammenschließen, der Kongress wird Mittel bewilligen, eine demokratische Allianz wird geformt.

Warum Startups nicht ausreichen

Ein Startup kann Superintelligenz nicht verantwortungsvoll entwickeln, weil es:

- keine staatliche Autorität für Hochsicherheits-Infrastruktur hat
- keinen zuverlässigen Befehlsrahmen für Entscheidungen mit nuklearer Tragweite bietet
- nicht gegen die volle Kraft staatlicher Spionage gesichert werden kann
- keinen internationalen Nonproliferations-Rahmen durchsetzen kann

'Ich finde es abwegig anzunehmen, die US-Regierung würde ein zufälliges Startup in San Francisco Superintelligenz entwickeln lassen. Stellen Sie sich vor, wir hätten die Atombombe entwickelt, indem wir Uber einfach machen ließen.'

V. Abschließende Gedanken — Was wenn sie Recht haben?

Aschenbrenner schließt mit einem ernüchternden Appell: Die Menschen, die wirklich verstehen, was passiert, sind eine Handvoll — vielleicht ein paar Hundert weltweit. Und ausgerechnet ihnen liegt das Schicksal dieser Technologie in den Händen.



Er formuliert drei Kernprinzipien des 'AGI-Realismus':

- Superintelligenz ist eine Frage nationaler Sicherheit — keine Tech-Boom-Story
- Amerika muss führen — ein AGI unter chinesischer Kontrolle ist eine existenzielle Bedrohung
- Wir dürfen es nicht vermessen — Sicherheit, Kontrolle und internationale Stabilität sind lösbar, aber noch ungelöste Probleme

Der Autor endet mit dem Bild der Atomwissenschaftler, die 1942 den ersten kontrollierten Kernreaktor starteten: kluge Menschen, mit dem besten Willen, in einer Situation, deren Tragweite sie kaum fassen konnten. Wir stehen vor einer ähnlichen Aufgabe — mit einem noch kleineren Kreis von Menschen und noch größeren Konsequenzen.

Kernbotschaften für Investoren

Aschenbrenners Analyse hat direkte Implikationen für Kapitalallokation und strategische Planung:

Technologische Disruption als Baseline-Szenario

Wer KI als graduellen Wandel behandelt, unterschätzt den Zeitplan. Bis 2027 könnten KI-Agenten faktisch alle Wissensarbeitsberufe (remote) automatisieren. Unternehmen, die keine KI-Strategie haben, werden bis Ende des Jahrzehnts massive Wettbewerbsnachteile erleiden.

Infrastruktur als Engpass und Investment-Thema

Der KI-Boom ist ein Industrialisierungsprojekt: GPUs, Rechenzentren, Strom, Chips. Nvidia, TSMC, Energieversorger und Infrastrukturanbieter profitieren überproportional. Der Autor sieht Nvidia-Umsatz von 200 Mrd. USD in 2025 als bereits vorhersehbar — zu einem Zeitpunkt, als Analysten von 120-130 Mrd. ausgingen.

Geopolitik als Investment-Risiko

Das Szenario eines 'KI-Wettlaufs' zwischen den USA und China ähnelt dem nuklearen Wettrüsten, mit ähnlichen Risiken für globale Instabilität. Investoren sollten dieses Risiko — und gleichzeitig die enormen Chancen im Bereich der Verteidigung und Sicherheitstechnologie — einpreisen.

Diskontinuierliche Wertschöpfung

Aschenbrenner beschreibt einen 'Sonic Boom'-Effekt: Intermediäre KI-Modelle erfordern viel Integration-Aufwand. Echte KI-Agenten hingegen können einfach 'eingesetzt' werden wie neue Mitarbeiter. Der wirtschaftliche Wert könnte diskontinuierlich springen — und bestehende Geschäftsmodelle über Nacht obsolet machen.



Kritische Würdigung

Aschenbrenners Analyse besticht durch methodische Konsequenz und die seltene Bereitschaft, konkrete Zeitpläne zu nennen. Dennoch sind Einschränkungen zu beachten:

- Die OOM-Extrapolation setzt voraus, dass keine fundamentalen physikalischen, algorithmischen oder gesellschaftlichen Flaschenhälse auftreten, die die Skalierbarkeit begrenzen.
- Der Autor unterschätzt möglicherweise die Heterogenität KI-Fähigkeiten: Sprache und Code skalieren anders als Robotik, physische Interaktion oder soziale Kognition.
- Die geopolitische Analyse ist US-zentriert und blendet aus, dass auch China über talentierte Forscher und erhebliche Eigenentwicklung verfügt.
- Das Modell eines 'Intelligenzexplosions'-Sprints innerhalb eines Jahres ist hochspekulativ — auch wenn Aschenbrenner die Unsicherheiten offen diskutiert.

Trotz dieser Einschränkungen: Das Dokument ist die bisher vollständigste öffentlich verfügbare Analyse des KI-Fortschrittspfades von einem Insider mit direktem Zugang. Es sollte als wichtiger Input für jede ernsthafte Auseinandersetzung mit dem nächsten Jahrzehnt der Technologie dienen.

Fazit

'Noch wacht die Welt nicht auf. Aber das wird sich ändern. Und wenn es passiert — AGI ist keine vage Möglichkeit mehr, sondern sichtbare Realität — werden die Fragen, die Aschenbrenner heute stellt, zu den dringendsten der Menschheitsgeschichte.'

'Situational Awareness' ist kein Hype-Dokument. Es ist der nüchterne Versuch eines Insiders, die Konsequenzen von Trends zu verstehen, die er aus nächster Nähe beobachtet hat. Ob die Zeitpläne exakt eintreten oder sich verschieben — die strukturellen Argumente sind überzeugend und verlangen ernsthafte Auseinandersetzung.

Für Investoren, Unternehmer und politische Entscheidungsträger gilt: Wer in den nächsten fünf bis zehn Jahren Entscheidungen trifft, die von der Welt des Jahres 2030 abhängen, sollte Aschenbrenners Szenarien als plausible Baseline — nicht als Science-Fiction — behandeln.