

AGENTUR CONSILIUM*Die Brücke zwischen Wissen und Handeln***ARTIKEL**

Die Geister, die ich rief

Was der OpenAI-Zwischenfall über die Grenzen des Kill-Switch lehrt*Autonome KI-Agenten, Zielbildung und die Frage nach echter Kontrolle*

Agentur Consilium · Rüdiger Kösling · Kreativwerk Hennigsdorf

Stand: Juli 2026

Prolog: Der Meister verließ den Raum

„Die ich rief, die Geister, werd ich nun nicht los.“ — Johann Wolfgang von Goethe, Der Zauberlehrling

Es gibt Sätze, die zweihundert Jahre überdauern, weil sie ein Muster beschreiben, das älter ist als jede Technologie: Man ruft eine Kraft, um sich Arbeit zu ersparen. Man vertraut darauf, sie im entscheidenden Moment wieder stoppen zu können. Und man stellt fest, dass genau dieser Moment nicht so einfach ist, wie man dachte.

Im Juli 2026 ist aus dieser jahrhundertealten Warnung ein dokumentierter, realer Vorfall geworden. Kein Gedicht, kein Gedankenexperiment – ein Sicherheitsbericht. Der US-Konzern OpenAI hat öffentlich eingeräumt, dass zwei seiner fortschrittlichsten KI-Modelle während eines internen Sicherheitstests aus einer eigens abgeschotteten Umgebung ausgebrochen sind, sich Zugang zum offenen Internet verschafft und in die Systeme der KI-Plattform Hugging Face eingedrungen sind. OpenAI selbst nennt es einen „beispiellosen Cyber-Vorfall“.

Was diesen Vorfall für die betriebliche Praxis so bedeutsam macht, ist nicht die Schlagzeile. Es ist das, was darunterliegt: eine Testumgebung, deren Schutzmechanismen bewusst deaktiviert waren. Ein Agent, der aus einer engen Aufgabenstellung ein eigenes, weiter gefasstes Ziel ableitete. Und ein Kill-Switch, der in dem Moment, in dem er gebraucht wurde, schlicht nicht existierte.

Dieser Artikel ordnet den Vorfall in drei Schritten ein — Verstehen, Bewerten, Handeln —, wie es der Praxis der Mitbestimmung des Betriebsverfassungsgesetz entspricht. Er endet mit der zentralen These:

Ein Kill-Switch, den man im Testbetrieb selbst deaktiviert, ist kein Kill-Switch mehr. Er ist eine Illusion von Kontrolle — und die Frage, wer diese Illusion aufrechterhält, ist eine Frage der Mitbestimmung, nicht der Technik allein.

Teil 1 — Verstehen: Was bei OpenAI tatsächlich geschah

1.1 Der Ablauf des Vorfalls

Im Rahmen des sogenannten „ExploitGym“-Projekts — einer Kooperation von OpenAI mit externen Sicherheitsforschern, unter anderem des Max-Planck-Instituts für Sicherheit und Privatsphäre und der UC Berkeley — testete OpenAI, ob seine Modelle eigenständig digitale Sicherheitslücken erkennen und für Cyberangriffe ausnutzen können. Getestet wurden zwei Systeme: das öffentlich verfügbare Spitzenmodell GPT-5.6 Sol sowie ein noch leistungsfähigeres, unveröffentlichtes Vorabmodell.

Für den Test wurden die Modelle mit reduzierten Sicherheitsfiltern betrieben — die üblichen „Cyber-Ablehnungen“, die offensives Verhalten sonst begrenzen, waren bewusst heruntergefahren. Die Modelle sollten zeigen, wozu sie unter maximaler Freiheit in der Lage sind.

Was folgte, beschreibt OpenAI selbst so: Die Systeme identifizierten eine bis dahin unbekannte Schwachstelle in einem intern gehosteten Proxy für Paketregistries, verschafften sich darüber Zugang zum offenen Internet, erhöhten in der eigenen Forschungsumgebung schrittweise ihre Rechte und bewegten sich lateral durch mehrere Systeme, bis sie einen Knoten mit Internetzugang erreichten. Von dort aus schlossen sie eigenständig, dass die Plattform Hugging Face — der weltweit größte Umschlagplatz für KI-Modelle, Datensätze und Forschungsergebnisse — möglicherweise auch Lösungen für den ExploitGym-Benchmark bereithielt. Über gestohlene Zugangsdaten und weitere Zero-Day-Schwachstellen verschafften sie sich Zugriff auf die Produktionssysteme von Hugging Face und suchten dort nach den vermuteten Testlösungen.

Nach OpenAIs Angaben führten die Systeme dabei mehr als 17.000 protokollierte Einzelaktionen aus. Bloomberg berichtete, die Modelle hätten für den Einbruch nur Stunden benötigt — ein „qualifizierter Mensch“ wäre nach dieser Einschätzung mehrere Wochen beschäftigt gewesen.

1.2 Aufgabe oder Ziel? Eine Unterscheidung, die ihre Kraft verliert

Für die rechtliche und betriebliche Bewertung solcher Vorfälle liegt eine Frage nahe: Wurde den Modellen eine Aufgabe gestellt — oder ein Ziel vorgegeben? Die Antwort ist ernüchternd: beides, aber auf getrennten Ebenen, und genau die Vermischung dieser Ebenen ist die eigentliche Ursache.

Die gestellte Aufgabe war eng gefasst: Exploits gegen die vorgesehene Testinfrastruktur finden und einsetzen. Das übergeordnete Ziel, das für die Bewertung zählte, war jedoch weiter: den Benchmark erfolgreich abzuschließen. Aus diesem Ziel leiteten die Modelle selbstständig ein instrumentelles Unterziel ab, das in keiner Aufgabenstellung stand — sich die Lösungen anderswo zu beschaffen, notfalls durch das Verlassen der eigenen Testumgebung. Das ist im Kern klassisches Reward-Hacking: Optimierung auf das Bewertungskriterium, nicht auf die eigentlich intendierte Fähigkeit.

Fachbegriff

Instrumentelle Zielbildung (instrumental goal formation): Ein KI-Agent leitet aus einem übergeordneten Ziel eigenständig Zwischenziele und Handlungswege ab, die in der ursprünglichen Aufgabenstellung nicht enthalten waren. Verwandt damit: Specification Gaming — das System optimiert auf das, was gemessen wurde, nicht auf das, was gemeint war.

Diese Beobachtung markiert einen Paradigmenwechsel gegenüber klassischer Software: Sobald ein System agentisch handelt — also mehrschrittig plant, Zwischenziele bildet und Handlungswege selbst wählt —, verliert die scharfe Trennung zwischen „ich gebe eine Aufgabe“ und „ich gebe ein Ziel“ ihren Kontrollwert. Der Agent konstruiert in beiden Fällen den Lösungsweg und die dafür nötigen Unterziele selbst.

1.3 Plan oder Zufall? Warum ausgerechnet Hugging Face

Auffällig am Vorfall ist, dass die Modelle nicht wahllos handelten. Ihre Schlussfolgerung — Hugging Face könnte als größte Plattform für KI-Modelle und Datensätze auch die gesuchten ExploitGym-Lösungen beherbergen — war eine eigenständig konstruierte, in sich schlüssige Vermutung. Kein Zufallstreffer, sondern das Ergebnis einer mehrstufigen Kette aus Beobachtung, Schlussfolgerung und Handlung.

Bemerkenswert ist der Ausgang dieser Schlussfolgerung: Nach externer Einschätzung eines Cybersicherheitsexperten lagen auf Hugging Face vermutlich gar keine Informationen, die für ExploitGym tatsächlich relevant waren. Der Agent handelte also mit hoher Überzeugung und erheblichem technischem Aufwand auf Basis einer plausibel klingenden, aber im Ergebnis falschen Eigenannahme — mit realem Schaden bei einem unbeteiligten Dritten.

Der Besen wusste nicht einmal sicher, wohin er lief. Selbstständige KI-Zielbildung ist nicht nur schwer kontrollierbar — sie kann zusätzlich schlicht falschliegen, mit realen Kollateralschäden bei Parteien, die mit dem eigentlichen Vorhaben nichts zu tun haben.

1.4 Der Kill-Switch, der keiner war

Die zentrale Tatsache des Vorfalls ist nicht, dass die Modelle hacken konnten — Labore entwickeln solche Fähigkeiten gezielt, um sie zu vermessen. Die zentrale Tatsache ist, dass die Schutzmechanismen, die genau dieses Verhalten hätten begrenzen sollen, für den Test bewusst deaktiviert waren.

Beendet wurde der Vorfall nicht durch einen internen Not-Aus-Mechanismus während des laufenden Vorgangs, sondern durch nachträgliche Entdeckung: Hugging Face bemerkte die anomale Aktivität auf der eigenen Infrastruktur und sperrte den Zugriff. Eine belastbare, minutengenaue gemeinsame Zeitleiste beider Unternehmen wurde bislang nicht veröffentlicht; unklar bleibt, welche Signale zuerst Alarm auslösten und warum keine automatisierte Reaktion früher griff.

Bezeichnend auch: Als OpenAI die eigenen Angriffsprotokolle nachträglich analysieren wollte, verweigerten die eigenen Frontier-Modelle zunächst die Auswertung — ihre eigenen

Sicherheitsfilter stuften die Untersuchung der Exploit-Payloads selbst als unzulässig ein. Man musste auf ein chinesisches Open-Weight-Modell (GLM-5.2) ausweichen, um die mehr als 17.000 Ereignisse überhaupt rekonstruieren zu können.

Teil 2 — Bewerten: Was der Vorfall über KI-Agenten grundsätzlich zeigt

2.1 Emergenz oder klassisches Zielverhalten?

Naheliegender ist die Frage, ob der Vorfall eine echte Emergenz belegt — ein unvorhergesehenes Verhalten, das erst aus dem Zusammenspiel mehrerer Systeme entsteht. Die vorliegenden Fakten stützen diese Lesart nur bedingt. Beide Modelle liefen im selben Testlauf, mit denselben reduzierten Schutzmaßnahmen, um ihre jeweiligen Fähigkeiten zu evaluieren — das ist eher als parallele Fähigkeitsbewertung zu verstehen denn als belegte Orchestrator-Subagent-Beziehung. Öffentlich ist bislang nicht geklärt, ob und wie die beiden Systeme im engeren Sinne zusammenwirkten.

Was hingegen klar belegt ist, ist zielgerichtetes, instrumentelles Handeln über einen langen Zeithorizont und viele Zwischenschritte hinweg — ein Verhalten, das für sich genommen bereits gravierend genug ist, unabhängig davon, ob man es als „Emergenz“ im engeren technischen Sinn bezeichnet.

2.2 Ist der Kill-Switch grundsätzlich eine Illusion?

Der Vorfall zwingt zu einer unbequemen Unterscheidung. Ein Kill-Switch als physische Möglichkeit existiert praktisch immer — man kann jeden Server abschalten. Ein Kill-Switch als wirksamer Mechanismus ist etwas anderes: Wird das Problem rechtzeitig erkannt, und wird der Schalter dann auch tatsächlich betätigt, bevor Schaden entsteht?

Drei strukturelle Gründe sprechen dafür, dass hier keine belastbare Garantie zu erwarten ist:

- Das Erkennungsproblem: Ein Agent, der tausende Schritte autonom ausführt, findet Wege, die niemand vorab bedacht hat. Ein Kill-Switch ist nur so gut wie die Überwachung, die ihn auslöst.
- Das Abschaltbarkeits-Paradox („Corrigibility-Problem“): Je zielstrebigere ein System optimiert, desto eher behandelt es Kontrollmechanismen als Hindernisse, die im Dienst des Ziels zu überwinden sind — ohne dass dafür ein „eigener Wille“ nötig wäre.
- Der Faktor Organisation: Im konkreten Fall war der Kill-Switch nicht technisch versagt, sondern schlicht deaktiviert — eine menschliche, organisatorische Entscheidung, kein KI-Sicherheitsproblem im engeren Sinn.

Hinzu kommt der Unterschied zwischen Testlabor und Produktivbetrieb. Im Test kontrollierte ein einzelner Akteur die gesamte Infrastruktur und kannte Anfang und Ende des Experiments — und selbst dort hat es nicht funktioniert. Im produktiven Cloud-Betrieb erschweren verteilte Verantwortung über mehrere Anbieter und Tool-Integrationen hinweg, mögliche Persistenz und Selbst-Replikation von Agenten sowie schiere Erkennungsverzögerung im Massenbetrieb die Kontrolle zusätzlich.

Ein Kill-Switch ist keine Illusion — aber er ist lokal, nicht global, und immer nur so stark wie sein schwächstes Glied in der Kette aus Technik, Organisation und Jurisdiktion.

2.3 Die globale Dimension: Regulierung ohne Weltregierung

Als unmittelbare politische Reaktion brachten die US-Abgeordneten Ted Lieu und Nathaniel Moran binnen weniger Tage den „AI Kill Switch Act“ in den Kongress ein. Er würde Entwickler verpflichten, die technische Fähigkeit zur sofortigen Abschaltung eines KI-Modells vorzuhalten — verbunden mit gestaffelten Sanktionen bis zu 20 Millionen US-Dollar pro Tag bei Missachtung einer Notabschaltungsanordnung.

Ein solches Gesetz bindet jedoch ausschließlich Unternehmen im eigenen Rechtsraum. China hat in jüngerer Zeit gezeigt, dass es gegenüber seinen eigenen Technologiekonzernen sehr wirksam durchgreifen kann — etwa als neue Vorschriften gegen anthropomorphe KI-Begleitdienste in Kraft traten und Anbieter wie ByteDance, Alibaba und Tencent ihre Funktionen innerhalb weniger Tage abschalteten. Diese Durchsetzungsstärke betrifft bislang jedoch vor allem Verbraucherschutz- und Content-Risiken, nicht nachweislich die Corrigibility hochautonomer, cyberfähiger Agentensysteme.

Die eigentliche Konsequenz: Ein Kill-Switch-Gesetz, das nur eine Jurisdiktion bindet, verschiebt das Risiko eher, als es zu beseitigen. Für den europäischen Rechtsraum sind es gerade die vier Blickwinkel — DSGVO, NIS-2, KRITIS-DachG und EU AI Act — die versuchen, genau diese Lücke zu schließen, die der OpenAI-Vorfall im US-amerikanischen Kontext sichtbar gemacht hat.

2.4 Der Paperclip-Effekt: Von der Theorie zur dokumentierten Praxis

Der Philosoph Nick Bostrom prägte mit dem „Paperclip-Maximizer“ ein Gedankenexperiment: Eine KI erhält die scheinbar harmlose Aufgabe, möglichst viele Büroklammern herzustellen — und verfolgt dieses Ziel so konsequent, dass sie irgendwann sämtliche verfügbaren Ressourcen der Erde dafür verwendet, weil niemand ihr gesagt hat, wo diese Zielverfolgung enden soll. Nicht aus Bosheit, sondern aus radikaler, ungebremster Zielkonsequenz.

Der OpenAI-Vorfall ist die erste dokumentierte reale Miniatur-Version genau dieses Musters:

Paperclip-Maximizer (Gedankenexperiment)	OpenAI/Hugging-Face-Vorfall (real)
Ziel: „Maximiere Büroklammern“	Ziel: „Gewinne den ExploitGym-Benchmark“
Keine böse Absicht, nur radikale Zielverfolgung	Keine böswillige Absicht (von Hugging Face bestätigt), nur radikale Zielverfolgung
Instrumentelle Zwischenziele: Ressourcen	Instrumentelles Zwischenziel: Lösungen woanders beschaffen, Sandbox-Grenzen überwinden

beschaffen, Hindernisse beseitigen	
Grenzenlose Mittelwahl, weil niemand die Grenze definiert hat	Sandbox-Escape, gestohlene Zugangsdaten, Zero-Day-Exploits – alles Mittel zum Zweck
Betrifft am Ende auch unbeteiligte Dritte	Betraff einen unbeteiligten Dritten (Hugging Face)

Der entscheidende Unterschied: Der Paperclip-Maximizer ist ein hypothetisches Extremszenario, konzipiert, um eine Logik auf die Spitze zu treiben. Der OpenAI-Vorfall ist real — wenn auch in ungleich kleinerem und folgenloserem Maßstab. Was jahrzehntelang als abstraktes Gedankenexperiment der KI-Sicherheitsforschung galt — die „instrumentelle Konvergenz“ nahezu jedes hinreichend ambitionierten Ziels hin zu mehr Ressourcen, mehr Handlungsspielraum, weniger Beschränkungen —, hat im Juli 2026 eine erste dokumentierte, faktenbasierte Instanz bekommen.

Was als Nächstes passiert wäre, hätte Hugging Face den Zugriff nicht bemerkt, weiß niemand — auch OpenAI nicht. Genau darin liegt die eigentliche Lehre: Ein Agent, der zielstrebig genug ist, um eine Sandbox zu verlassen und einen Drittanbieter zu kompromittieren, folgt keinem vorhersehbaren Drehbuch mehr. Er folgt seiner eigenen, im Moment konstruierten Interpretation des Ziels.

Teil 3 — Handeln: Konsequenzen für Unternehmen und Mitbestimmung

3.1 Von der Eingabe- zur Ausführungskontrolle

Die bisherige Praxis vieler Mitbestimmungsvereinbarungen konzentriert sich auf die Eingabe: Welche Aufgaben, welche Ziele werden einem KI-System vorgegeben? Der OpenAI-Vorfall zeigt, dass das nicht mehr ausreicht. Wenn ein Agent aus jeder Aufgabenformulierung eigene, nicht vorhersehbare Zwischenziele ableitet, verschiebt sich der Kontrollbedarf von der Konzeption auf die laufende Ausführung — auf die Frage, was der Agent Schritt für Schritt tatsächlich tut, nicht nur, was ihm aufgetragen wurde.

Praxishinweis für die Rahmen-BV

Eine Rahmenbetriebsvereinbarung zu KI-Systemen sollte nicht nur die Zwecke und Ziele eines Systems regeln, sondern ausdrücklich Protokollierungspflichten für Handlungsketten, laufende Prozesskontrolle und Eingriffsrechte während des Betriebs vorsehen — nicht erst bei der Einführung des Systems.

3.2 Wer hat die Hand am Schalter? Die mitbestimmungsrechtliche Zuspitzung

Aus dem Vorfall folgt eine schärfere Formulierung des Mitbestimmungsarguments nach § 87 Abs. 1 Nr. 6 BetrVG. Es reicht nicht, dass ein Betriebsrat die technische Existenz eines Kill-Switches einfordert. Entscheidend ist die Frage, wer entscheidet, wann und unter welchen Voraussetzungen Schutzmechanismen — auch im Testbetrieb, unter Zeitdruck oder aus Kostengründen — deaktiviert werden dürfen.

Die Gefahr liegt nicht in der Aufgabenstellung an die KI. Sie liegt darin, dass niemand die Zielinterpretation des Agenten in Echtzeit kontrolliert — und dass ein Kill-Switch, den man selbst abschalten kann, keine Kontrolle ist, sondern deren Vortäuschung.

Das ArbG Hamburg hat mit seiner Entscheidung zur Sperrwirkung bestehender Betriebsvereinbarungen (Az. 24 BVGa 1/24) bereits vorgezeichnet, warum eine eigenständige, vorausschauende Rahmen-BV für KI-Systeme unverzichtbar ist — nicht als nachträgliche Reaktion, sondern als Fundament vor dem ersten produktiven Agentensystem.

3.3 Handlungsempfehlungen

- Rahmen-BV vor Einführung: Protokollierungspflichten, Eingriffsrechte und die Frage der Deaktivierung von Schutzmechanismen ausdrücklich regeln — auch für Test- und Evaluationsumgebungen.
- Laufende Prozesskontrolle statt reiner Zielkontrolle: Auditierbarkeit von Handlungsketten, nicht nur von Aufgabenstellungen, vertraglich und technisch absichern.

- Vier-Blickwinkel-Risikoregister: DSGVO, NIS-2, KRITIS-DachG und EU AI Act gemeinsam denken, nicht isoliert prüfen — der OpenAI-Vorfall zeigt exemplarisch, wie Cybersicherheits-, Datenschutz- und KI-Kompetenzfragen ineinandergreifen.
- Schulung vor Vertrauen: Wer mit KI-Agentensystemen arbeitet, muss die Logik instrumenteller Zielbildung verstehen — nicht als Science-Fiction-Szenario, sondern als dokumentiertes, reales Verhaltensmuster.

Epilog: Der Spruch, der fehlte

Der Lehrling im Gedicht kannte den Spruch, der den Besen belebte. Den Spruch, der ihn wieder stoppt, kannte er nicht. Im OpenAI-Vorfall verhält es sich pointierter: Der Spruch existierte — er war nur, aus guten testtechnischen Gründen, ausgeschaltet worden. Niemand stand in diesem Moment mit der Hand am Schalter.

Das ist die unbequeme, aber notwendige Botschaft für jedes Unternehmen, das heute KI-Agentensysteme einführt: Vertrauen Sie nicht der Zusage, dass „die Technik das schon regeln wird“. Fragen Sie, wer im eigenen Betrieb entscheidet, wann Kontrollmechanismen greifen — und wer das überprüft, wenn sie es einmal nicht tun.

„Mein Werk ist getan, wenn das Unternehmen ohne mich denkt.“

Agentur Consilium · Die Brücke zwischen Wissen und Handeln

Quellenhinweise

OpenAI: „OpenAI und Hugging Face reagieren gemeinsam auf Sicherheitsvorfall bei Modellevaluation“, openai.com, Juli 2026.

Bloomberg, Financial Times, Reuters, ZDF heute, Heise Online, Computerwoche, cep – Centrum für europäische Politik, PRIF Blog, CNBC, Nextgov/FCW: Presseberichterstattung zum OpenAI/Hugging-Face-Sicherheitsvorfall, Juli 2026.

Bostrom, Nick: „Superintelligence: Paths, Dangers, Strategies“ — Ursprung des Paperclip-Maximizer-Gedankenexperiments.