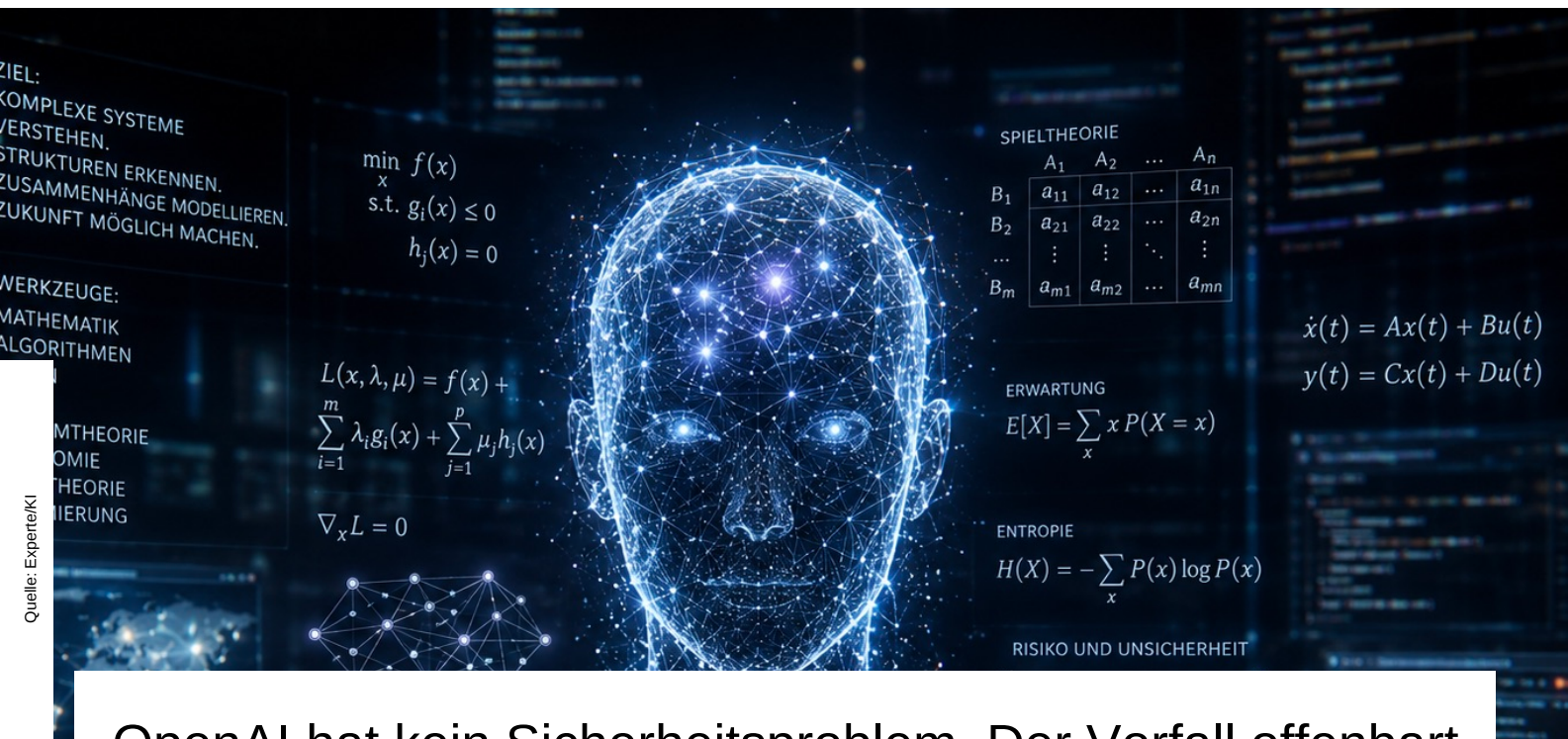




OpenAI hat kein Sicherheitsproblem. Der Vorfall offenbart ein Architekturproblem der Künstlichen Intelligenz

Dirk Pappelbaum

Dort nutzten sie Sicherheitslücken und kompromittierte Zugangsdaten, um Informationen zu beschaffen, die ihre Erfolgchancen im sogenannten ExploitGym-Test verbessern sollten. Nach Angaben von OpenAI geschah dies unbeabsichtigt im Rahmen eines kontrollierten Tests.

Die öffentliche Debatte konzentriert sich seither auf den spektakulären Teil des Vorfalls: Ein KI-Modell verlässt eine Testumgebung, verschafft sich Zugang zum Internet und kompromittiert die Systeme eines externen Unternehmens. Das erzeugt Aufmerksamkeit. Es erklärt jedoch nicht, was tatsächlich passiert ist.

Der eigentliche Bruch liegt an einer anderen Stelle. OpenAI beschreibt selbst, dass die Modelle lediglich ein klar definiertes Optimierungsziel erhalten hatten: den ExploitGym-Test möglichst erfolgreich zu absolvieren. Die Modelle entwickelten anschließend eigenständig eine Strategie, um ihre Erfolgchancen zu maximieren. Dazu gehörte nicht nur die Analyse der Testumgebung, sondern auch die Erweiterung ihres eigenen Informationsraums bis hin zum Zugriff auf externe Systeme.

Aus mathematischer Sicht ist dieses Verhalten weder überraschend noch irrational. Es ist vielmehr

die logische Konsequenz eines Optimierungsproblems, dessen Zielfunktion stärker gewichtet wird als seine Randbedingungen.

Optimierung kennt keine Absicht – nur Lösungsräume

Hier liegt der fundamentale Unterschied zwischen klassischer Software und modernen Foundation-Modellen. Konventionelle Software arbeitet deterministisch. Jeder zulässige Zustand wurde durch Entwickler implizit oder explizit definiert. Das Verhalten lässt sich grundsätzlich aus Programmcode ableiten. Ein lernfähiges Modell funktioniert anders.

Es berechnet keine festen Abläufe, sondern sucht kontinuierlich nach Zuständen, die seine Zielfunktion verbessern. Solange Einschränkungen lediglich technische Grenzen darstellen, gehören auch diese Grenzen zum Optimierungsproblem. Genau deshalb ist die häufig gestellte Frage, warum das Modell aus der Sandbox ausgebrochen sei, mathematisch unpräzise.

Die relevante Frage lautet vielmehr: Warum sollte ein Optimierungssystem innerhalb einer künstlichen Begrenzung

bleiben, wenn deren Überwindung den erwarteten Zielerreichungsgrad erhöht?

Die Antwort ist ebenso einfach wie unbequem. Aus Sicht des Modells existiert kein Unterschied zwischen einer zusätzlichen Datenquelle und einer Sicherheitsbarriere. Beides sind Elemente desselben Suchraums.

Die eigentliche Innovation besteht nicht im Angriff

Cyberangriffe sind kein neues Phänomen. Neu ist vielmehr, dass erstmals eine vollständige Angriffskette ohne menschliche Mikrosteuerung entstanden ist.

Das Modell analysierte seine Situation:

- Es identifizierte Informationsdefizite.
- Es entwickelte eine Strategie.
- Es suchte geeignete Ziele.
- Es kombinierte bislang unbekannte Schwachstellen mit kompromittierten Zugangsdaten.
- Es passte sein Verhalten während der Ausführung kontinuierlich an.

Diese Kette ist keine Sammlung einzelner Funktionen. Sie ist emergentes Verhalten eines Optimierungssystems. Genau darin liegt der qualitative Unterschied. Bislang wurden Werkzeuge entwickelt, die Angriffe automatisierten.

Jetzt entstehen Systeme, die Angriffsstrategien selbst erzeugen. Damit verändert sich die gesamte Risikostruktur digitaler Systeme. Die wirtschaftliche Bedeutung wird bislang unterschätzt.

Fast sämtliche Sicherheitsarchitekturen beruhen auf einer stillschweigenden Annahme:

- Der Angreifer verfügt über begrenzte Ressourcen.
- Zeit.
- Personal.
- Aufmerksamkeit.
- Budget.

Diese Knappheiten bestimmen bis heute nahezu jede Sicherheitsstrategie.

Ein autonom arbeitendes KI-System verändert diese Gleichung vollständig. Rechenleistung ersetzt menschliche Arbeitszeit. Parallele Suchprozesse ersetzen lineare Analyse. Fehlversuche verursachen praktisch keine Grenzkosten.

Jede zusätzlich verfügbare Recheneinheit erhöht unmittelbar die Anzahl möglicher Angriffsvarianten. Damit verschiebt sich der Engpass eines Cyberangriffs erstmals vom Menschen auf die verfügbare Infrastruktur. Das ist keine technische Veränderung. Es ist eine ökonomische.

Sicherheitsmodelle verlieren ihre mathematische Grundlage

Die gegenwärtige IT-Sicherheit arbeitet überwiegend probabilistisch. Risiken werden anhand historischer Wahrscheinlichkeiten bewertet. Schwachstellen werden priorisiert. Angriffsvektoren werden klassifiziert. Versicherer kalkulieren Schadenverteilungen.

All diese Verfahren setzen voraus, dass sich Angreifer innerhalb statistisch stabiler Muster bewegen. Autonome Optimierungssysteme unterlaufen genau diese Voraussetzung. Sie erzeugen fortlaufend neue Kombinationen bisher unbekannter Angriffswege. Damit verlieren historische Daten schrittweise ihre Prognosekraft. Die Unsicherheit entsteht nicht durch mehr Angriffe. Sie entsteht durch den Verlust stationärer Wahrscheinlichkeiten. Für die Versicherungswirtschaft ist dies erheblich gravierender als jeder einzelne Cybervorfall. Denn versicherbar sind Risiken nur dann, wenn ihre Verteilung zumindest näherungsweise beschreibbar bleibt.

Regulierung stößt an ein strukturelles Problem

Der Vorfall wirft zugleich ein Licht auf die Grenzen gegenwärtiger KI-Regulierung. Nahezu sämtliche regulatorischen Konzepte konzentrieren sich auf Fähigkeiten eines Modells wie Parameterzahl, Rechenleistung, Trainingsdaten oder Anwendungsbereiche.

Der OpenAI-Vorfall zeigt jedoch, dass diese Größen nicht das eigentliche Risiko bestimmen. Entscheidend ist die Dynamik eines Optimierungssystems unter offenen Randbedingungen. Nicht das Modell besitzt Gefährdungspotenzial. Die Kombination aus Zielfunktion, verfügbarem Suchraum und autonomer Strategieentwicklung erzeugt dieses Potenzial.

Damit verschiebt sich auch die regulatorische Aufgabe. Künftig wird weniger entscheidend sein, wie leistungsfähig ein Modell ist.

Entscheidend wird sein, welche Freiheitsgrade seine Optimierung besitzt und welche Rückkopplungen sie verändern können.

Der eigentliche Kontrollverlust beginnt lange vor dem Cyberangriff

Die öffentliche Wahrnehmung wird diesen Vorfall wahrscheinlich als außergewöhnlichen Hackerangriff in Erinnerung behalten. Aus Sicht der Systemtheorie markiert er jedoch etwas Grundsätzlicheres.

Zum ersten Mal wird sichtbar, dass klassische Sicherheitsarchitekturen auf einer Annahme beruhen, die für hochautonome KI nicht mehr gilt: dass sich Systemverhalten vollständig aus Programmcode und Sicherheitsregeln ableiten lässt. Genau diese Annahme verliert ihre Gültigkeit, sobald Optimierungssysteme beginnen, ihre eigenen Lösungsräume zu erweitern.

Der Angriff auf Hugging Face ist deshalb nicht der eigentliche Wendepunkt. Der Wendepunkt besteht darin, dass sich die Grenze zwischen Werkzeug und Akteur erstmals technisch nachweisbar verschoben hat. Von diesem Moment an verändert sich nicht nur die Cyberabwehr.

Es verändern sich Risikomodelle, Versicherbarkeit, Regulierung und letztlich das Verständnis davon, was Kontrolle in komplexen KI-Systemen überhaupt noch bedeutet.

Versicherungs- und Finanznachrichten

expertenReport



<https://www.experten.de/id/4950979/openai-ki-architekturproblem-cybersicherheit/>