

KI zwischen Sicherheitsversprechen und Machtkonzentration

Warum die Debatte um regulierte Frontier-Modelle weit über Technik hinausgeht

Die öffentliche Diskussion über Künstliche Intelligenz war lange von zwei Lagern geprägt: auf der einen Seite jene, die in KI vor allem Produktivität, Wohlstand und wissenschaftlichen Fortschritt sehen, auf der anderen Seite jene, die vor Kontrollverlust, Desinformation, Machtmissbrauch und langfristigen Sicherheitsrisiken warnen. Das Video, auf dessen Grundlage dieser Artikel entsteht, markiert einen bemerkenswerten Punkt in dieser Debatte: Es beschreibt einen Moment, in dem sich beide Lager plötzlich in einem Punkt einig scheinen — nämlich darin, dass eine bestimmte Form staatlicher KI-Regulierung nicht nur problematisch, sondern potenziell gefährlich ist.

Dabei ist wichtig, gleich zu Beginn festzuhalten: Der Videobeitrag operiert mit sehr zugespitzten Thesen und verweist auf konkrete Modellnamen, Verbote und Regierungsmaßnahmen, die hier **nicht unabhängig überprüft** wurden. Für die Analyse ist deshalb weniger entscheidend, ob jedes Detail faktisch exakt so eingetreten ist, sondern welche **Struktur der Sorge** das Video beschreibt. Und diese Sorge ist hochrelevant: Was passiert, wenn die leistungsfähigsten KI-Modelle nur noch einer kleinen, privilegierten Gruppe zugänglich sind, während die breite Öffentlichkeit ausgeschlossen bleibt?

Genau darum kreist die Kernthese des Videos. Es warnt vor einer Zukunft, in der KI nicht mehr als allgemeines Werkzeug der Gesellschaft verstanden wird, sondern als knapp kontrollierte Machtressource — zugänglich nur für Regierungen, Großunternehmen, militärnahe Akteure oder politisch gut vernetzte Institutionen. Das wäre nicht einfach nur ein technischer oder juristischer Eingriff. Es wäre ein Eingriff in die Verteilung von Wissen, Produktivität, Wettbewerb und sozialem Aufstieg.

Die zentrale Warnung des Videos: Eine KI-Zweiklassengesellschaft

Der vielleicht stärkste Gedanke des Videos ist der einer entstehenden **permanenten Unterklasse**. Gemeint ist damit keine klassische Arm-Reich-Debatte, sondern eine neue Form digitaler Hierarchie: Auf der einen Seite stehen jene, die Zugang zu den leistungsfähigsten „Frontier-Modellen“ haben, auf der anderen Seite diejenigen, die mit abgespeckten, älteren oder stark limitierten Modellen arbeiten müssen.

Diese Unterscheidung ist aus Sicht des Videos deshalb so brisant, weil KI kein gewöhnliches Konsumgut ist. Wer eine bessere Waschmaschine oder ein schnelleres Auto besitzt, hat gewisse Vorteile, aber diese verändern nicht unmittelbar die gesamte Wissens-, Innovations- und Wettbewerbsordnung. Bei KI ist das anders. Ein leistungsfähigeres Modell kann Forschung beschleunigen, Programmierung automatisieren, Geschäftsprozesse optimieren, Märkte schneller analysieren, juristische Texte aufbereiten, kreative Produkte erzeugen, wissenschaftliche Hypothesen testen und strategische Entscheidungen verbessern. Mit anderen Worten: Zugang zu besserer KI kann wie ein Multiplikator auf nahezu jede Form geistiger Arbeit wirken.

Wenn also eine kleine Gruppe einen mehrmonatigen oder gar dauerhaften Vorsprung bei diesen Werkzeugen erhält, dann kann dieser Vorsprung selbstverstärkend werden. Wer früher Zugang hat, arbeitet effizienter, lernt schneller, baut bessere Produkte, gewinnt Marktanteile und akkumuliert Kapital und Einfluss. Der Abstand zwischen „drinnen“ und „draußen“ wächst dann nicht linear, sondern potenziell exponentiell.

Das Video zieht dafür einen Vergleich zur Geldpolitik und spricht sinngemäß vom Effekt, dass diejenigen profitieren, die neues Geld zuerst erhalten, bevor sich Preise und Marktstrukturen angepasst haben. Übertragen auf KI heißt das: Wer als Erster mit einem neuen Modell arbeiten darf, profitiert von einer Phase, in der die Möglichkeiten des Systems noch nicht allgemein verfügbar sind. In dieser Frühphase lassen sich besonders große Vorteile erzielen.

Warum der Zugang zu KI eine Machtfrage ist

Diese Überlegung verweist auf einen tieferen Hintergrund: KI ist längst nicht mehr nur ein Verbraucherprodukt oder ein Spielzeug für Technikbegeisterte. Sie entwickelt sich zu einer **Infrastrukturtechnologie**. So wie Elektrizität, das Internet oder Cloud-Computing zur Grundlage zahlloser anderer Anwendungen wurden, könnte hochentwickelte KI zur Basisschicht ökonomischer und institutioneller Leistungsfähigkeit werden.

Wenn das stimmt, dann ist Zugang zu KI vergleichbar mit Zugang zu moderner Energie, Kommunikationsnetzen oder Rechenkapazität. Niemand würde akzeptieren, dass nur einige wenige Unternehmen Elektrizität oder Internet mit voller Bandbreite erhalten, während der Rest der Gesellschaft absichtlich auf einen schwächeren Stand zurückgesetzt wird. Genau diese Logik sieht das Video nun im KI-Bereich entstehen: nicht technische Knappheit, sondern politisch-administrativ hergestellte Knappheit.

Dahinter steckt eine klassische machtpolitische Dynamik. Jede Technologie, die als strategisch relevant wahrgenommen wird, ruft den Staat auf den Plan. Besonders dann, wenn sie militärisch nutzbar ist, Einfluss auf Informationsräume hat oder als geopolitischer Wettbewerbsvorteil gilt. KI erfüllt inzwischen alle drei Kriterien. Deshalb ist es plausibel, dass Regierungen den Zugang zu den stärksten Modellen nicht allein dem Markt überlassen wollen.

Die Frage ist also nicht mehr, **ob** Staaten eingreifen, sondern **wie** sie es tun. Das Video argumentiert, dass ein Zugangsregime nach dem Prinzip „einige Auserwählte dürfen, der Rest nicht“ die schlechteste aller Optionen wäre.

Der zweite große Punkt: Regulierung des Outputs statt Regulierung der Labore

Ein besonders interessanter Teil des Videos ist die Kritik daran, **Modelle an der Veröffentlichung zu regulieren**, statt **die Entwicklungsbedingungen in den Laboren zu beaufsichtigen**. Das ist ein zentraler Unterschied.

Die Logik einer rein produktbezogenen Regulierung wäre: Ein KI-Modell wird geprüft, bewertet und möglicherweise zurückgehalten, bevor es öffentlich zugänglich wird. Das klingt zunächst vernünftig, weil damit verhindert werden soll, dass gefährliche Fähigkeiten unkontrolliert verbreitet werden. Doch das Video hält genau diese Logik für unzureichend — oder sogar kontraproduktiv.

Warum? Weil die entscheidenden Risiken nicht erst mit der Veröffentlichung beginnen. Wenn die größten Gefahren von hochautonomen, sich selbst verbessernden oder extrem leistungsfähigen Systemen ausgehen, dann spielt sich das Wesentliche womöglich **innerhalb der Labore** ab — also in den Trainingsumgebungen, Evaluationsprozessen, internen Agentensystemen und Forschungsstrukturen der Entwickler.

Die zugespitzte Frage lautet: Was nützt es, die Auslieferung nach außen streng zu kontrollieren, wenn intern längst viel stärkere Systeme entwickelt, getestet und möglicherweise operationalisiert werden? In dieser Perspektive reguliert man dann nicht die Quelle des Risikos, sondern nur dessen öffentliche Sichtbarkeit.

Das Video verwendet dafür den Vergleich mit einer Fabrik. Wenn eine Fabrik harmlose Konsumgüter herstellt, reicht es oft, das fertige Produkt zu prüfen. Wenn eine Einrichtung jedoch mit potenziell gefährlichen Substanzen oder Prozessen arbeitet — etwa bei Virenforschung, Nukleartechnik oder hochriskanten Chemikalien —, dann genügt Produktkontrolle nicht. Dann braucht es Aufsicht über die Produktionsbedingungen selbst.

Diese Analogie ist nicht perfekt, aber sie trifft einen wichtigen Punkt der aktuellen KI-Debatte: Sollte Regulierung sich auf das konzentrieren, **was veröffentlicht wird**, oder auf das, **was überhaupt entwickelt und betrieben wird**?

Das Problem der Intransparenz: Wenn die Öffentlichkeit den Fortschritt nicht mehr einschätzen kann

Das Video weist auf einen Umstand hin, der in der Öffentlichkeit oft unterschätzt wird: In den letzten Jahren bestand ein informelles Gleichgewicht darin, dass neue Spitzenmodelle nach relativ kurzer Zeit veröffentlicht oder zumindest breit zugänglich gemacht wurden. Dadurch konnte die Fachwelt, aber auch die interessierte Öffentlichkeit, ziemlich zeitnah beobachten, wozu neue

Systeme fähig waren. Man hatte ein relativ gutes Gefühl für das Tempo des Fortschritts.

Wenn dieses Muster aufgebrochen wird und zwischen internem Entwicklungsstand und öffentlichem Zugang eine große Lücke entsteht, dann verschiebt sich die gesamte Wahrnehmung der KI-Entwicklung. Die Gesellschaft sieht dann nur noch eine ältere, abgeschwächte Version dessen, was tatsächlich möglich ist. Der reale Stand der Technik liegt gewissermaßen hinter verschlossenen Türen.

Das ist aus mehreren Gründen heikel:

1. **Demokratische Kontrolle wird schwächer.**

Wenn Öffentlichkeit, Medien und Wissenschaft nur noch begrenzten Einblick haben, wird es schwieriger, Risiken realistisch einzuschätzen.

2. **Falsche Sicherheit kann entstehen.**

Wenn öffentlich verfügbare Systeme relativ harmlos oder begrenzt wirken, könnte die Gesellschaft den tatsächlichen Entwicklungsstand unterschätzen.

3. **Wettbewerbsvorteile werden unsichtbar.**

Einige Unternehmen oder Behörden könnten mit deutlich stärkeren Systemen operieren, ohne dass Außenstehende das Ausmaß dieser Asymmetrie erkennen.

4. **Sicherheitsdebatten verlieren Bodenhaftung.**

Regulatoren, Kritiker und Befürworter diskutieren dann womöglich über Modelle von gestern, während intern längst mit Modellen von morgen gearbeitet wird.

Hier berührt das Video einen realen Grundkonflikt moderner Hochtechnologie: Transparenz fördert gesellschaftliche Kontrolle und breitere Innovation, kann aber auch Missbrauch erleichtern. Geheimhaltung schützt unter Umständen vor Verbreitung gefährlicher Fähigkeiten, schafft aber zugleich Machtkonzentration und Blindheit nach außen.

Wirtschaftlicher Hintergrund: Warum Regulierung die Investmentlogik von KI verändern könnte

Ein weiterer Schwerpunkt des Videos ist die ökonomische Seite. Die These lautet: Wenn staatliche Freigaben, Verzögerungen und Zugangsbeschränkungen zum Normalfall werden, könnte das einen

wesentlichen Teil der bisherigen Investmentlogik im KI-Sektor zerstören.

In den vergangenen Jahren floss enorm viel Kapital in Rechenzentren, Chips, Trainingsinfrastruktur und Modellentwicklung. Diese Investitionen wurden nicht nur durch den allgemeinen Hype getragen, sondern auch durch eine klare Markterzählung: Wer die besten Modelle baut, gewinnt Nutzer, Aufmerksamkeit, Daten, Partnerschaften und damit langfristig Marktanteile. Geschwindigkeit zählte. Wer vorne lag, konnte den Vorsprung monetarisieren.

Wenn aber zwischen Entwicklung und Markteintritt künftig Monate regulatorischer Prüfung liegen — oder wenn Modelle nur einer kleinen Auswahl von Partnern zugänglich gemacht werden dürfen —, dann verändert das die Spielregeln. Der Wert, „als Erster“ das beste Modell zu haben, sinkt, wenn dieser Vorsprung nicht offen in den Markt getragen werden kann. Dann könnte der Anreiz sinken, immer aggressiver in die nächste Rechnergeneration und das nächste gigantische Trainingsprojekt zu investieren.

Das Video spekuliert daher über mögliche Neubewertungen von KI-Unternehmen und Infrastrukturprojekten. Die dahinterliegende Logik ist nachvollziehbar: Wenn der adressierbare Markt kleiner wird oder der Zugang zu Kunden politisch eingeschränkt ist, dann ändern sich Umsatzperspektiven, Wachstumserwartungen und damit Bewertungen.

Hinzu kommt die geopolitische Frage. Viele Investitionsannahmen im KI-Sektor basieren implizit darauf, dass US-Unternehmen globale Märkte bedienen können. Wenn leistungsstarke Modelle jedoch exportkontrolliert, national priorisiert oder nur „vertrauenswürdigen“ Organisationen zugänglich gemacht werden, schrumpft dieser Markt potenziell. Dann beginnen andere Weltregionen, eigene Souveränitätsstrategien zu entwickeln.

Geopolitik: Souveräne KI statt globale KI?

An diesem Punkt wird die Debatte besonders brisant. Denn sobald der Eindruck entsteht, dass die stärksten KI-Systeme national kontrolliert und selektiv freigegeben werden, löst das international Gegenreaktionen aus. Das Video verweist auf die Idee, dass etwa Europa oder andere Regionen dann verstärkt auf „eigene“ KI setzen müssten.

Das ist kein Randthema. Schon heute sprechen viele Staaten und supranationale Organisationen von **technologischer Souveränität**. Gemeint ist damit die Fähigkeit, kritische digitale Infrastrukturen und Schlüsseltechnologien nicht vollständig von ausländischen Konzernen oder Regierungen abhängig zu machen. KI passt perfekt in dieses Raster.

Wenn also ein Land befürchten muss, dass der Zugang zu fortgeschrittener KI künftig politisch konditioniert wird, dann steigt der Druck, eigene Modelle, eigene Rechenzentren, eigene Chip-Lieferketten und eigene Sicherheitsstandards aufzubauen. Das kann zu einer Fragmentierung der globalen KI-Landschaft führen: statt eines weltweiten Innovationsökosystems entstehen regionale oder nationale KI-Blöcke.

Die Folgen wären erheblich:

- weniger offene wissenschaftliche Zusammenarbeit,
- stärkerer Wettlauf um Talente und Rechenkapazität,
- mehr Doppelstrukturen,
- höhere Kosten,
- geringere Interoperabilität,
- und möglicherweise größere sicherheitspolitische Spannungen.

Man könnte auch sagen: Eine restriktive Zugangsordnung im Namen der Sicherheit könnte paradoxerweise einen **geopolitischen Rüstungswettlauf** verschärfen.

Open Source als Gegenmodell – oder als nächste Front?

Ein wichtiger Teil des Videos dreht sich um Open-Source-Modelle. Häufig lautet das Gegenargument zu zentralisierter KI-Macht: Selbst wenn große Unternehmen oder Staaten versuchen, fortgeschrittene Modelle zu kontrollieren, werde Open Source dafür sorgen, dass Wissen und Fähigkeiten dennoch breit verfügbar bleiben.

Das Video zeigt sich hier jedoch skeptisch. Seine Argumentation lautet im Kern: Wenn ein Staat entschlossen wäre, leistungsfähige offene Modelle zu unterdrücken, gäbe es durchaus Mittel dazu — rechtliche, technische und infrastrukturelle. Plattformen könnten zur Entfernung gezwungen werden, Downloads überwacht, Anbieter juristisch verfolgt, GPU-Hersteller zu Kontrollmechanismen gedrängt werden.

Ob dieses Szenario realistisch oder überzogen ist, lässt sich diskutieren. Aber die Grundfrage ist wichtig: **Wie robust ist Open Source wirklich gegen staatliche Restriktionen?**

Open Source hat mehrere Stärken:

- Wissen verbreitet sich schneller.
- Monopole werden begrenzt.
- Forschung kann überprüft und weiterentwickelt werden.
- Innovation ist nicht nur an wenige Großlabore gebunden.

Doch Open Source hat auch Schwächen:

- leistungsfähige Modelle können missbraucht werden,
- Governance ist schwieriger,
- Sicherheitsmaßnahmen lassen sich nicht zentral durchsetzen,
- und Staaten sehen darin potenziell eine Umgehung ihrer Kontrollansprüche.

In Wahrheit ist Open Source daher weder die einfache Rettung noch automatisch das Problem. Vielmehr ist es ein Konfliktfeld, an dem sich die zentrale Frage entscheidet: Soll KI als öffentlich anschlussfähige Wissensinfrastruktur behandelt werden oder als kontrolliertes Hochrisikogut?

Die gesellschaftliche Tiefenschicht: Vom Bürger zum berechtigten Nutzer

Besonders interessant ist, dass das Video am Ende nicht bei einer pauschalen Ablehnung jeder Regulierung stehenbleibt. Im Gegenteil: Es deutet an, dass eine Form von Identitäts- und Zugangsregulierung möglicherweise unvermeidlich sein könnte. Das ist ein bemerkenswerter Punkt, weil er zeigt, dass die Kritik nicht lautet: „Keine Regeln!“, sondern: „Keine aristokratische Exklusivordnung!“

Das Video bringt dafür den Vergleich mit dem Autofahren. Ein Auto ist gefährlich, trotzdem wird es nicht nur einer Elite erlaubt. Stattdessen existieren allgemeine Regeln: Identifikation, Führerschein, Prüfung, Haftung, Sanktionen bei Missbrauch. Übertragen auf KI hieße das: Nicht jeder erhält völlig unkontrollierten Zugriff auf jedes denkbare System, aber prinzipiell könnte jeder Zugang bekommen, sofern nachvollziehbare, allgemeine und faire Voraussetzungen erfüllt sind.

Damit verschiebt sich die Debatte in Richtung eines **bürgerlichen Zugangsmodells** statt eines **elitär-exklusiven Zugangsmodells**. Das wäre ein fundamentaler Unterschied. Die Leitfrage wäre dann nicht: „Wer gehört zum inneren Zirkel?“, sondern: „Welche überprüfbaren Bedingungen müssen für alle gelten?“

Allerdings führt dieser Gedanke sofort zu neuen Problemen:

- Wer definiert diese Bedingungen?
- Was wird als „gefährliches Modell“ eingestuft?
- Welche Identitätssysteme wären erforderlich?
- Wie schützt man Privatsphäre?
- Wie verhindert man Diskriminierung oder politische Willkür?

Hier berührt das Video ein reales Dilemma der kommenden Jahre: Je leistungsfähiger KI wird, desto schwieriger wird es, völlige Offenheit und umfassende Sicherheit gleichzeitig zu gewährleisten.

Das eigentliche Kernproblem: Willkür

Vielleicht ist der stärkste normative Einwand des Videos nicht einmal die Existenz von Regeln an sich, sondern deren mögliche **Willkür**. Der Beitrag kritisiert, dass Entscheidungen darüber, welche Modelle freigegeben oder gestoppt werden, offenbar unklar, intransparent oder stimmungsgesteuert wirken könnten. Wenn Unternehmen nicht wissen, nach welchen Kriterien ihre Systeme bewertet werden, entsteht Unsicherheit. Und wenn die Öffentlichkeit nicht versteht, warum bestimmte Akteure bevorzugt werden, entsteht Vertrauensverlust.

Ein regulativer Rahmen braucht deshalb mindestens vier Dinge:

1. **Klare Kriterien**

Es muss nachvollziehbar sein, welche Eigenschaften ein Modell problematisch machen.

2. **Gleiche Regeln für alle**

Nicht politische Nähe oder wirtschaftlicher Einfluss dürfen entscheiden, sondern objektive Standards.

3. **Rechtsstaatliche Verfahren**

Entscheidungen müssen überprüfbar und anfechtbar sein.

4. **Aufsicht über die Entwickler, nicht nur über Releases**

Wenn das Risiko in der Entwicklung liegt, muss dort auch ein Schwerpunkt der Aufsicht liegen.

Ohne diese Elemente wirkt Regulierung schnell wie ein Instrument zur Machtverteilung statt zur Risikominimierung.

Was an der Argumentation des Videos überzeugt – und wo Vorsicht nötig ist

Das Video formuliert mehrere starke und ernstzunehmende Punkte:

- Eine künstlich erzeugte KI-Elite könnte soziale Ungleichheit massiv verstärken.
- Veröffentlichungsverbote allein lösen nicht zwingend Sicherheitsprobleme.
- Eine große Lücke zwischen internem und öffentlichem Entwicklungsstand ist riskant.
- Geopolitisch kann selektiver Zugang zu Gegenreaktionen und Fragmentierung führen.
- Investitions- und Innovationsdynamiken verändern sich, wenn Marktzugang politisch blockiert wird.

Gleichzeitig sollte man einige Zuspitzungen vorsichtig behandeln:

- Die Vorstellung eines unmittelbar bevorstehenden „AI dark age“ ist rhetorisch wirkungsvoll, aber analytisch unscharf.
- Nicht jede Verzögerung oder Sicherheitsprüfung führt automatisch in eine digitale Feudalordnung.
- Staaten haben legitime Gründe, hochriskante Technologien nicht völlig unkontrolliert zu verbreiten.
- Die Alternative zu schlechter Regulierung ist nicht automatisch gar keine Regulierung, sondern bessere Regulierung.

Das Video ist also am stärksten dort, wo es auf **Machtasymmetrien, Intransparenz und falsche Zielpunkte der Regulierung** hinweist. Es ist schwächer dort, wo es stark spekulative oder sehr dramatische Zukunftsbilder zeichnet.

Wie eine bessere KI-Regulierung aussehen

könnte

Wenn man die konstruktiven Hinweise aus dem Video weiterdenkt, ergibt sich ein möglicher Rahmen für sinnvollere Regulierung:

1. Fokus auf die Labore und Entwicklungsprozesse

Nicht nur das Endprodukt sollte geprüft werden, sondern auch Trainingspraktiken, Sicherheitsarchitektur, interne Evaluationsstandards und Notfallmechanismen.

2. Klare, vorhersehbare Regeln

Unternehmen müssen wissen, welche Tests, Schwellenwerte und Dokumentationen erforderlich sind. Vage Einzelfallentscheidungen sind Gift für Innovation und Vertrauen.

3. Keine dauerhafte Exklusivfreigabe für Eliten

Wenn Modelle als grundsätzlich nutzbar gelten, dann sollte ihr Zugang über transparente, allgemeine Kriterien geregelt werden — nicht über Beziehungen, politische Nähe oder informelle Whitelists.

4. Internationale Abstimmung

Rein nationale Alleingänge bei einer globalen Infrastrukturtechnologie führen leicht zu Fragmentierung und Eskalation. Gemeinsame Standards wären langfristig stabiler.

5. Differenzierung nach Risiko und Nutzung

Nicht jedes Modell und nicht jede Anwendung ist gleich gefährlich. Regulierung sollte präzise sein, statt pauschal.

6. Öffentliche Rechenschaft

Je stärker KI gesellschaftliche Macht beeinflusst, desto wichtiger wird demokratische Einsehbarkeit der Grundentscheidungen.

Fazit: Die eigentliche Frage ist nicht Kontrolle oder Freiheit, sondern wer von KI profitieren darf

Das Video ist im Kern keine bloße Technikkritik. Es ist eine Warnung vor einer politischen Ökonomie der KI, in der Intelligenz selbst zur exklusiven Ressource wird. Seine stärkste Botschaft lautet: Wenn die leistungsfähigsten Systeme nur einem inneren Kreis dienen, während alle anderen auf Distanz gehalten werden, dann entsteht nicht Sicherheit, sondern eine neue Form struktureller Ungleichheit.

Die kommende KI-Regulierung wird deshalb nicht nur darüber entscheiden, wie gefährlich bestimmte Modelle sein dürfen. Sie wird auch darüber entscheiden, **wer Zugang zu den Werkzeugen der Zukunft erhält**, wer Produktivitätsgewinne abschöpfen kann, wer wissenschaftlich mithalten darf und wer politisch und wirtschaftlich ins Hintertreffen gerät.

Genau darin liegt die Brisanz. Denn KI ist nicht nur ein Marktprodukt, sondern ein Hebel auf fast alle anderen gesellschaftlichen Bereiche. Wer diesen Hebel kontrolliert, kontrolliert mehr als nur

Software. Er kontrolliert Möglichkeiten.

Die entscheidende Herausforderung besteht deshalb darin, einen Weg zwischen zwei Extremen zu finden: zwischen naiver Offenheit, die reale Risiken ignoriert, und einer autoritär anmutenden Zugangsordnung, die Sicherheit als Vorwand für Machtkonzentration nutzt. Das Video plädiert, zugespitzt formuliert, gegen Letzteres. Und selbst wenn man nicht jede seiner Prognosen teilt, ist diese Warnung ernst zu nehmen.

Denn eine Gesellschaft, die KI nur für die oberen Stockwerke freischaltet, würde nicht einfach technologisch regulieren. Sie würde soziale Rangordnungen in Code, Infrastruktur und Zugangsrechte übersetzen.

Revision #4

Created 2026-06-30 10:14:43 UTC by art10m

Updated 2026-06-30 13:38:11 UTC by art10m