

Digitaler Kannibalismus: Frisst die generative KI das Web auf, das sie füttert?

Das Paradoxon des Parasiten: Eine Welt, die von einem unsichtbaren Wirt ernährt wird

Wir stehen am Beginn einer neuen Ära, die von künstlicher Intelligenz angetrieben wird. Im Zentrum dieses Wandels steht ein faszinierendes Gedankenexperiment: Wenn jeder nur noch KI-Anwendungen wie Google Gemini nutzt, um Antworten zu erhalten, welchen Wert hat dann die Arbeit der Millionen von Menschen, die das Wissen, die Kreativität und die Daten als Grundlage dafür quasi kostenfrei zur Verfügung stellen? Ist das nicht zutiefst unfair? Diese Frage ist keine akademische Spitzfindigkeit. Sie deckt ein grundlegendes Paradoxon auf, das im Herzen des aktuellen KI-Ökosystems schlägt: Die generative KI funktioniert in ihrer jetzigen Form wie ein Parasit, der sich von einem riesigen, unsichtbaren Wirt – den von Menschen geschaffenen Inhalten des Internets – ernährt, während sie gleichzeitig die ökonomischen Grundlagen untergräbt, die diesen Wirt am Leben erhalten.

Das Ausmaß der Nahrungsaufnahme

Um die Dimensionen dieses Prozesses zu verstehen, muss man sich die schiere Menge an Daten vor Augen führen, die für das Training eines großen Sprachmodells (Large Language Model, LLM) wie Gemini erforderlich ist. Es handelt sich nicht um ein abstraktes „Lernen“, sondern um das systematische, groß angelegte Kopieren und Verarbeiten der gesamten Bandbreite menschlichen digitalen Ausdrucks. LLMs werden mit einer „massiven Datenmenge“ trainiert, die Milliarden von Texten und anderen Inhalten umfasst.[1] Die Grundlage dafür bilden riesige Web-Crawls von gemeinnützigen Organisationen wie **Common Crawl**. Diese Organisation archiviert Petabytes an Daten und fügt monatlich 3 bis 5 Milliarden neue Seiten hinzu.[2, 3] Diese Datensätze sind das Fundament, das durch kuratierte Quellen wie Wikipedia, Bücher (sowohl gemeinfreie als auch illegal beschaffte), wissenschaftliche Arbeiten und Code-Archive ergänzt wird.[4, 5, 6] OpenAI gibt explizit an, öffentlich zugängliche Informationen aus dem Internet, Daten von Drittpartnern und Nutzerinformationen zu verwenden.[7]

Diese Abhängigkeit ist nicht nur fundamental, sondern birgt auch ein systemisches Risiko. LLMs benötigen einen konstanten Strom an riesigen, vielfältigen und aktuellen Daten, um relevant und präzise zu bleiben.[1, 5, 8] Diese Daten werden überwiegend von menschlichen Kreativen produziert, die innerhalb verschiedener Wirtschaftsmodelle agieren.[9, 10, 11] Das aktuelle KI-Modell stört jedoch genau diese Modelle, indem es den Traffic und damit die Einnahmen reduziert.[12, 13] Indem die KI-Industrie ihre eigene Datenquelle untergräbt, schafft sie eine langfristige Nachhaltigkeitskrise für sich selbst. Wenn der „Wirt“ geschwächt wird, wird der „Parasit“ letztendlich unter einem Mangel an frischer, hochwertiger Nahrung leiden, was zu einer Stagnation oder Verschlechterung der Modelle führen wird.

Die wahre Natur der „Bestie“: Stochastische Papageien, keine denkenden Maschinen

Um die ethischen und rechtlichen Implikationen zu bewerten, ist es entscheidend, die Funktionsweise eines LLM zu verstehen. Es handelt sich nicht um ein bewusstes Wesen, das „versteht“,

sondern um ein hochkomplexes statistisches Modell. Der Begriff „stochastischer Papagei“ beschreibt dies treffend: Es sind Systeme, die statistische Zusammenhänge nutzen, um überzeugend menschenähnlichen Text zu generieren, ohne dabei ein echtes semantisches Verständnis zu besitzen.[14, 15] Sie sind darauf trainiert, das wahrscheinlichste nächste Wort in einer Sequenz vorherzusagen.[4, 8, 16] Sie sind Meister der Nachahmung und Mustererkennung, nicht des Verstehens. Obwohl sie emergente Fähigkeiten zur logischen Schlussfolgerung zeigen können, ist ihre Grundlage probabilistisch, nicht kognitiv.[17] Diese Unterscheidung ist von zentraler Bedeutung, da sie die später zu diskutierende Analogie zum „menschlichen Lernen“ entkräftet.

Die Ungerechtigkeit im Kern

Hier liegt das zentrale ethische Dilemma. Der Wert wird von Millionen von Menschen geschaffen, aber von einer Handvoll Technologiegiganten erfasst und kommerzialisiert, die für das grundlegende Rohmaterial nicht bezahlt haben. Der Prozess umfasst das Scraping von allem, von persönlichen Blogs und Nachrichtenartikeln bis hin zu urheberrechtlich geschützten Büchern aus Pirateriequellen wie Books3.[18, 19, 20] Dies geschieht, um kommerzielle Produkte zu entwickeln, die, wie der Anstieg der Marktkapitalisierung von Microsoft um eine Billion Dollar zeigt, immens wertvoll sind.[21]

Die KI-Industrie verwendet den Begriff „Trainingsdaten“, um dieses Material zu beschreiben. Dies ist jedoch ein beschönigender Euphemismus, der die wahre Natur der verwendeten Inhalte verschleiert. Der Begriff „Daten“ impliziert rohe, unstrukturierte und wertneutrale Informationen.[1, 5, 8] Tatsächlich handelt es sich jedoch um „Werke“ – Romane, Artikel, Fotografien, Code –, von denen jedes ein fertiges Produkt menschlicher Arbeit, Kreativität und geistigen Eigentums darstellt.[18, 21, 22] Indem „kreative Werke“ als „Trainingsdaten“ umgedeutet werden, minimiert die Industrie den wahrgenommenen Wert und die geistigen Eigentumsansprüche der ursprünglichen Schöpfer. Der Akt des massenhaften Kopierens erscheint so eher als neutraler technischer Prozess denn als die Aneignung fertiger Güter.

1 Der alte Pakt ist gebrochen: Wie die KI den Klick getötet hat

Jahrzehntlang basierte das digitale Veröffentlichungswesen auf einem ungeschriebenen Pakt zwischen Inhaltserstellern und Suchmaschinen. Dieser implizite Vertrag besagte: Google und andere durchsuchen und indexieren Inhalte und schicken im Gegenzug wertvollen, monetarisierbaren Traffic an die Websites der Ersteller. Dieser Pakt wurde nun einseitig von der KI aufgekündigt, was eine existenzielle Krise für die gesamte digitale Medienlandschaft auslöst.

1.1 Der Motor der Creator Economy

Die Creator Economy, von einzelnen Bloggern bis hin zu großen Nachrichtenorganisationen, hat ihre Arbeit historisch durch eine Vielzahl von Einnahmequellen finanziert, die fast alle vom Website-Traffic abhängig sind. Dazu gehören Werbeeinnahmen (oft an Seitenaufrufe oder Klicks gekoppelt), Affiliate-Marketing, gesponserte Beiträge, Abonnements und Paywalls, der Verkauf digitaler Produkte und direkte Spenden.[9, 10, 11, 23, 24, 25, 26, 27, 28, 29] Über Jahrzehnte war die Google-Suche der primäre Motor, der den Traffic lieferte, der diese Modelle antrieb.

1.2 Die große Kannibalisierung

Die Einführung von KI-gestützten Zusammenfassungen (AI Overviews) und Chatbots hat diese Dynamik radikal verändert. Indem sie direkte Antworten liefern, machen sie es für den Nutzer überflüssig, auf die ursprüngliche Quell-Website zu klicken. Dadurch wird die lebenswichtige Verbindung zwischen Inhaltserstellung und Monetarisierung gekappt. Die Daten zeichnen ein düsteres Bild:

- Verleger berichten von „plötzlichen und anhaltenden“ Traffic-Rückgängen von 25 % bis 30 %.[12]
- Der Eigentümer der Daily Mail gibt an, dass die Klickraten bei einigen Suchanfragen um bis zu 89 % gesunken sind.[12]
- Eine Studie unter den 500 meistbesuchten Publishern zeigt einen Rückgang der Besuche um 27 % im Jahresvergleich.[30]
- Die World History Encyclopedia verzeichnete einen Einbruch des Traffics um 25 %, nachdem sie in den AI Overviews prominent platziert wurde.[31]

Führungskräfte aus der Medienbranche beschreiben dies als einen „zweigleisigen Angriff“ und eine „existenzielle Krise“.[12] Der Wandel von der „Suche“ zur „Antwort“ gestaltet die ökonomische Geografie des Internets grundlegend um. Das alte Web war ein Netzwerk von Zielen (Websites), und Suchmaschinen waren die Verzeichnisse, die die Reise dorthin erleichterten. Der Wert (Werbeeinnahmen, Abonnements) wurde am Zielort erfasst. KI-Schnittstellen wie Gemini und ChatGPT werden nun selbst zum Ziel. Sie verweisen nicht nur auf Informationen, sie synthetisieren und präsentieren sie. Da die Nutzer ihre Antworten direkt von der KI erhalten, schwindet die Notwendigkeit, die ursprünglichen Quell-Websites zu besuchen.[30, 31] Dies führt zu einer massiven Verlagerung der wirtschaftlichen Macht. Der Wert, der einst auf Millionen von Websites verteilt war, wird nun innerhalb der KI-Plattformen konsolidiert, die über ihnen stehen.

1.3 Das gebrochene Versprechen

Die offiziellen Erklärungen von Google, dass die KI-Suche „qualitativ hochwertige Klicks“ fördere, stehen in direktem Widerspruch zu den Daten und Erfahrungen der Verleger. Liz Reid, Leiterin der Google-Suche, behauptet, die KI treibe „mehr Suchanfragen und qualitativ hochwertige Klicks“ an.[12] Publisher-Verbände wie Digital Content Next (DCN) und die Professional Publishers Association (PPA) legen jedoch Daten vor, die Rückgänge von 10 % bis 25 % im Jahresvergleich belegen.[13] Der implizite Vertrag – Google crawlt Inhalte und schickt im Gegenzug Traffic – ist zerbrochen.[31]

Dieser Wandel schafft auch ein „Vertrauensdefizit“, das nicht nur den Verlagen, sondern auch dem Nutzen der KI selbst schadet. Früher verlieh ein hohes Ranking bei Google selbst unbekannten Publishern eine gewisse Glaubwürdigkeit, da die Quelle sichtbar war.[13] AI Overviews synthetisieren Informationen oft ohne prominente, klare Quellenangabe oder, schlimmer noch, „halluzinieren“ Fakten und schreiben sie fälschlicherweise seriösen Quellen zu.[32, 33] Dies untergräbt das Vertrauen der Nutzer in die bereitgestellten Informationen. Wenn Nutzer die Quelle nicht leicht überprüfen können oder der Zusammenfassung der KI nicht trauen, sinkt der Wert der „Antwort“. Dies führt zu einem Teufelskreis: Verlage verlieren Traffic und den Anreiz, qualitativ hochwertige, überprüfbare Inhalte zu produzieren, was wiederum die Trainingsdaten für zukünftige KIs verschlechtert und sie noch unzuverlässiger macht.

2 Der Raubzug am helllichten Tag: Urheberrecht, „Fair Use“ und die falsche Analogie des „Lernens“

Die Frage der „Unfairness“ findet ihren juristischen und ethischen Kern im Konflikt um das Urheberrecht. Die Verteidigung der KI-Unternehmen stützt sich auf die Doktrin des „Fair Use“ (in den USA) und eine verführerische, aber zutiefst fehlerhafte Analogie: die Gleichsetzung des maschinellen Trainings mit menschlichem Lernen. Diese Argumentation ist ein eigennütziger Versuch, eine beispiellose Urheberrechtsverletzung zu legitimieren.

2.1 Das juristische Schlachtfeld

Mehrere wegweisende Klagen definieren derzeit die Grenzen dieses Konflikts:

- **The New York Times gegen OpenAI & Microsoft:** In dieser Klage wird massiver Urheberrechtsmissbrauch vorgeworfen. Es wird dargelegt, wie ChatGPT nahezu wörtliche Auszüge aus Artikeln hinter der Paywall reproduziert und damit direkt mit dem Kernprodukt der Times konkurriert und dieses entwertet.[21, 32, 34, 35]
- **Authors Guild gegen Anthropic:** Diese Klage konzentriert sich auf die Verwendung von raubkopierten Buchdatensätzen. Ein historischer Vergleich in Höhe von 1,5 Milliarden US-Dollar erkennt den Schaden an, der durch die Verwendung illegal beschaffter Werke entsteht.[19, 20, 36]
- **Getty Images gegen Stability AI:** Hier geht es um das Scraping von Millionen urheberrechtlich geschützter Bilder, einschließlich der Reproduktion von Wasserzeichen in KI-generierten Ausgaben, was auch eine Markenrechtsverletzung darstellt.[22, 37, 38, 39]

2.2 Die Dekonstruktion von „Fair Use“

Die „Fair Use“-Doktrin des US-Urheberrechts basiert auf vier Faktoren: (1) Zweck und Charakter der Nutzung, (2) Art des urheberrechtlich geschützten Werks, (3) Umfang und Wesentlichkeit des verwendeten Teils und (4) Auswirkung auf den Markt.[40, 41] KI-Unternehmen argumentieren, ihre Nutzung sei „transformativ“, da das Ziel darin bestehe, ein Modell zu trainieren und nicht, das ursprüngliche Werk neu zu veröffentlichen.[40, 42] Sie behaupten, der Trainingsprozess sei analog zum Lesen von Büchern durch einen Menschen zu Forschungszwecken.[43, 44]

Die „Fair Use“-Verteidigung ist jedoch ein strategischer Versuch, ein Geschäftsmodell nachträglich zu legalisieren, das auf der Annahme beruhte, dass eine massenhafte Urheberrechtsverletzung zu schwierig oder zu teuer zu verfolgen wäre. Der Aufbau eines grundlegenden Modells erfordert die Aufnahme des größtmöglichen Datensatzes.[1, 5] Eine Lizenzierung dieser Inhalte vor dem Training wäre unerschwinglich teuer und logistisch unmöglich gewesen. Die Industrie verfolgte daher einen Ansatz des „erst scrapen, dann um Verzeihung bitten“ und nahm alles, was öffentlich zugänglich war (und einiges, was es nicht war, wie raubkopierte Bücher).[18, 19] Das „Fair Use“-Argument ist kein Prinzip, mit dem sie begannen, sondern ein juristischer Schutzschild, den sie nun im Nachhinein zu errichten versuchen, um ihr Multi-Billionen-Dollar-Unternehmen zu schützen. Das Aufkommen eines Lizenzmarktes [45] untergräbt den vierten Faktor des Fair Use – den Marktschaden – tödlich, indem es beweist, dass ein Markt existiert, der geschädigt werden kann.

2.3 Die falsche Analogie: Warum KI-Training NICHT menschliches Lernen ist

Die zentrale Verteidigung der KI-Industrie bricht bei genauerer Betrachtung zusammen.

1. **Maßstab und Perfektion:** Menschen lernen aus einer begrenzten Anzahl von Werken und behalten „unvollkommene Eindrücke“ zurück.[42, 46] Das KI-Training beinhaltet die Erstellung perfekter, vollständiger digitaler Kopien von Millionen oder Milliarden von Werken in einem übermenschlichen Maßstab.[46, 47] Ein Student kann nicht legal die gesamte Bibliothek kopieren, um zu „lernen“.[48]
2. **Kommerzielle Ausbeutung:** Menschliches Lernen dient der persönlichen Bereicherung oder der Schaffung eines neuen, eigenständigen Werks. Das KI-Training ist ein direkter, zwischengeschalteter Schritt beim Aufbau eines kommerziellen Produkts, das oft direkt mit den Quellen konkurriert, auf denen es trainiert wurde.[48, 49] Der Gesellschaftsvertrag, der davon ausgeht, dass Menschen aus veröffentlichten Werken lernen, hat nie eine Maschine

vorweggenommen, die dies tut, um ein konkurrierendes Produkt in großem Maßstab zu schaffen.[46]

3. **Fehlende Kognition:** Die Analogie des „Lernens“ ist eine gefährliche Anthropomorphisierung.[46, 50, 51] KI-Systeme sind keine Menschen; sie sind komplexe Werkzeuge. Sie „verstehen“ nicht und haben keine „Ideen“. Sie replizieren statistisch Muster aus ihren Trainingsdaten.[14] Diesen mechanischen Prozess mit menschlicher Kognition gleichzusetzen, ist ein Kategorienfehler, der dazu dient, die Schutzmechanismen und Freiheiten zu beanspruchen, die wir menschlichem Denken gewähren.

Die juristischen Ergebnisse schaffen zudem einen zersplitterten und widersprüchlichen Präzedenzfall, bei dem die *Methode* der Beschaffung härter beurteilt wird als der *Akt* des Trainings selbst. Im Fall Anthropic entschied der Richter, dass das Training mit *legal erworbenen* Werken Fair Use sei, die Beschaffung dieser Werke durch Piraterie jedoch nicht.[43, 44, 52] Dies schafft eine bizarre Rechtslandschaft. Es legt nahe, dass der Akt des massenhaften Kopierens für das Training zulässig ist, solange die ursprüngliche Kopie legal beschafft wurde (z. B. durch den Kauf eines Buches vor dem Scannen). Diese Logik adressiert jedoch nicht den Kernschaden für die Schöpfer: die Schaffung eines abgeleiteten, konkurrierenden Produkts, das das Original entwertet, unabhängig davon, wie die Trainingsdaten beschafft wurden.

3 Eine Gabelung am Weg: Die Gestaltung einer neuen digitalen Wirtschaft

Nach der Analyse des Problems ist es an der Zeit, die aufkommenden Lösungen zu untersuchen. Die Auseinandersetzung zwischen KI-Entwicklern und Inhaltserstellern erzwingt die Entwicklung neuer ökonomischer Modelle. Es zeichnen sich verschiedene Wege ab, die die Zukunft der Inhaltenmonetarisierung und -vergütung im Zeitalter der KI prägen könnten.

3.1 Der Pakt der Barone: Pauschale Lizenzgeschäfte

Der vorherrschende Trend sind massive, hinter verschlossenen Türen ausgehandelte Abkommen zwischen den KI-Giganten und den größten Medienkonzernen. OpenAI hat Verträge mit Verlagen wie Associated Press, Axel Springer, Le Monde und der Financial Times unterzeichnet, deren Wert von 1-5 Millionen US-Dollar pro Jahr bis zu den 250 Millionen US-Dollar über fünf Jahre für News Corp reicht.[53, 54, 55, 56, 57] Diese Deals bieten den KI-Firmen Rechtssicherheit und den großen Verlagen eine neue Einnahmequelle.

3.2 Die Revolution der Mikrozahlungen: Nutzungsbasierte Vergütung

Es entstehen auch granularere und potenziell gerechtere Modelle, die die Bezahlung an die tatsächliche Nutzung knüpfen. Perplexity AI hat ein Modell zur Umsatzbeteiligung vorgeschlagen, bei dem Verlage einen „zweistelligen“ Prozentsatz der Werbeeinnahmen erhalten, wenn ihre Inhalte zitiert werden.[58] Bill Gross von ProRata schlägt ein Modell vor, das eine „Crawl-Gebühr“ (z. B. einen Cent pro Crawl) mit einer 50/50-Lizenzgebühr kombiniert, wenn Inhalte in einer Antwort verwendet werden, ähnlich dem Spotify-Modell.[59]

3.3 Der Fonds der Schöpfer: Direkte Vergütung für das Training

Einige Plattformen gehen einen anderen Weg und bezahlen die Kreativen direkt für ihren Beitrag zu den Trainingsdatensätzen selbst. Adobe hat einen Bonuspool für seine Adobe-Stock-Beitragenden eingerichtet, deren Arbeiten zum Training seines Firefly-Modells verwendet werden.[60] Canva hat einen 200-Millionen-Dollar-Fonds geschaffen, um Kreative zu entschädigen, die zustimmen, dass ihre Arbeiten für das KI-Training genutzt werden.[60]

3.4 Der souveräne Verleger: Geschlossene Gärten und markeneigene LLMs

Eine defensive Strategie für Verlage besteht darin, ihre eigenen proprietären KI-Tools zu entwickeln. Verlage können ihre Inhalte lizenzieren, um ihre eigenen, markengeschützten Chatbots zu trainieren, die auf KI-Infrastruktur (wie der von OpenAI) laufen, aber ausschließlich auf ihren eigenen Archiven trainiert sind.[59] Der Policy Intelligence Assistant von Politico Pro ist ein frühes Beispiel dafür. Er schafft ein hochwertiges, differenziertes Produkt, das seine einzigartigen Daten monetarisiert.[59]

Die folgende Tabelle fasst diese aufkommenden Modelle zusammen und vergleicht ihre Mechanismen und potenziellen Auswirkungen.

Tabelle 1: Ein Vergleich aufkommender KI-Vergütungsmodelle

Modell	Funktionsweise	Hauptbefürworter/Mögliche Beispiele	Mögliche Vorteile	Mögliche Nachteile
Pauschale Lizenzierung	KI-Firmen zahlen hohe, pauschale Gebühren für den Zugriff auf ganze Inhaltsarchive.	OpenAI & News Corp [55], Microsoft & Axel Springer [55]	Garantierte Einnahmen für Verlage; rechtliche Klarheit für KI-Firmen.	Bevorzugt große Akteure; schließt unabhängige Kreative aus; fehlt nutzungsbasierter Wert; intransparente Bedingungen.
Umsatzbeteiligung	KI-Plattformen teilen einen Prozentsatz der Werbeeinnahmen, die neben KI-Antworten generiert werden.	Perplexity AI [58]	Verknüpft Vergütung direkt mit Nutzung/Sichtbarkeit; schafft eine fortlaufende Partnerschaft.	Komplex zu verfolgen und zu prüfen; Einnahmen anfangs möglicherweise gering; abhängig vom Werbeerfolg der KI-Plattform.
Crawl-Gebühren & Lizenzgebühren	Mikrozahlungen für jeden gecrawlten Inhalt, plus Lizenzgebühren bei Verwendung in einer Ausgabe.	ProRata (Bill Gross) [59]	Granular und theoretisch fair; bewertet jeden Inhalt; skalierbar.	Technisch anspruchsvoll in der Umsetzung; erfordert universelle Akzeptanz; Potenzial für geringe Gesamtauszahlungen.
Direkte Vergütungsfonds	Plattformen legen einen Fonds auf, um Boni an Kreative zu zahlen, deren Arbeit im Training verwendet wird.	Adobe Firefly [60], Canva [60]	Direkte Bezahlung an Kreative; erkennt den Wert der Trainingsdaten selbst an.	Auszahlungen oft willkürlich, intransparent und nicht an die tatsächliche Wirkung gekoppelt; kann als symbolische Geste angesehen werden.

Fortsetzung auf nächster Seite

Tabelle 1 – Fortsetzung von vorheriger Seite

Modell	Funktionsweise	Hauptbefürworter/Mögliche Beispiele	Mögliche Vorteile	Mögliche Nachteile
Markeneigene LLMs	Verlage nutzen KI-Technologie, um proprietäre Tools zu entwickeln, die ausschließlich auf ihren eigenen Inhalten trainiert sind.	Politico Pro [59]	Volle Kontrolle über Inhalt und Monetarisierung; schafft ein einzigartiges Premium-Produkt; baut einen defensiven Schutzwall auf.	Hohe Entwicklungskosten; nur für Marken mit starken, umfangreichen und einzigartigen Archiven realisierbar.

4 Fazit: Das menschliche Signal bewerten, bevor es im Rauschen untergeht

Wir kehren zur Ausgangsfrage zurück: Welchen Wert hat die menschliche Arbeit in einer Welt, die von KI-Antworten dominiert wird? Die Analyse zeigt, dass das aktuelle Modell nicht nur unfair, sondern auch langfristig selbstzerstörerisch ist. Wir riskieren, in einen informationellen Ouroboros zu geraten – eine Endlosschleife, in der KIs auf den synthetischen, abgeleiteten Ausgaben anderer KIs trainiert werden. Dies würde zu einem geschlossenen Kreislauf von immer faderen, fehleranfälligeren und homogenisierten Inhalten führen, ohne neue menschliche Erfahrungen, Untersuchungen und Kreativität.

Der eigentliche Wert von von Menschen geschaffenen Inhalten liegt nicht in ihrer Nützlichkeit als statistisches Muster für eine Maschine. Ihr Wert liegt in ihrer Verbindung zur Realität: gelebte Erfahrung, verifizierte Fakten, originäre Forschung und einzigartiger künstlerischer Ausdruck. Das ist das „menschliche Signal“.

Es ist an der Zeit für einen neuen, expliziten digitalen Pakt. Dieser Rahmen muss über die zerbrochene implizite Vereinbarung der Suchmaschinen-Ära hinausgehen. Er erfordert eine Kombination aus technologischen Lösungen (wie transparente Quellenangaben und Nutzungsverfolgung), rechtlicher Klarheit (Reform des Urheberrechts oder Klärung von Fair Use für KI) und neuen Wirtschaftsmodellen (wie die in Abschnitt 4 untersuchten), die die menschliche Schöpfung angemessen bewerten und vergüten. Das Ziel ist nicht, die KI aufzuhalten, sondern sie in ein Ökosystem zu integrieren, das weiterhin Anreize für die menschliche Arbeit schafft, die ihr überhaupt erst einen Wert verleiht. Ein Scheitern wäre nicht nur unfair gegenüber den Schöpfern – es würde letztendlich die KI selbst wertlos machen.