

Künstliche Intelligenz und Öffentlichkeit

Felix Stalder

23–28 minutes

KI Slop ist überall. Übernimmt KI die öffentliche Sphäre? Bestimmt sie den Diskurs? Die einfache Antwort auf diese Frage ist: Nein! KI ist nicht und wird nicht autonom, aber Menschen nutzen diese Technologien, um damit den öffentlichen Diskurs in ihrem Sinne zu prägen.

Eine etwas differenzierte Antwort muss aber lauten: KI schreibt mit! Diese Erkenntnis ist nicht neu, sondern nimmt eine Erfahrung von Friedrich Nietzsche mit seiner schlecht funktionierenden Schreibmaschine auf, der 1882 in einem Brief feststellte: „[Unser Schreibzeug arbeitet mit an unseren Gedanken.](#)“

Technologien, so lässt sich Nietzsches oft zitierter Satz verstehen, verändern die Voraussetzungen des Denkens (und des Handelns). Marshall McLuhans berühmte Formel dafür lautete: „the medium is the message“. Medien transportieren nicht nur Ideen und Inhalte, sondern prägen grundsätzlich die Bedingungen ihrer Artikulation. Dies wird in Zeiten medialen Wandels besonders sichtbar: Das galt für Nietzsche, als er von der Feder zur Schreibmaschine zu wechseln versuchte, und es gilt für uns, die wir uns mit neuen Formen des Medienwandels befassen müssen. Damit stellt sich die Frage, wie KI am öffentlichen Diskurs mitschreibt – und wie wir diese Art und Weise des Mitschreibens beeinflussen können.

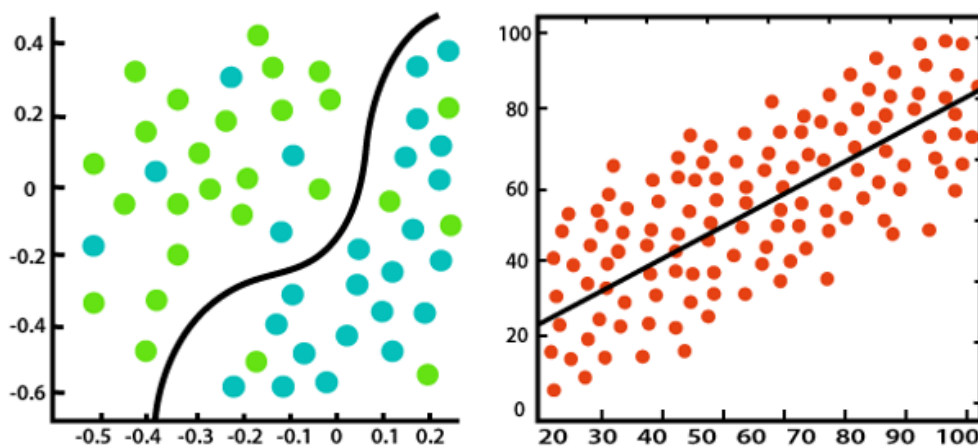
Klassifikation und Regression. Grundoperationen der KI

Beginnen wir mit der einfachsten Frage: Wie funktioniert die aktuell dominante Form der KI? Technisch gesehen berechnet sie die Wahrscheinlichkeit einer Annahme („auf diesem Bild ist eine

Katze“), auf Basis eines statistischen Verhältnisses einer Vielzahl von meist unbekannten Variablen, die aus einem oft unüberschaubar großen Datensatz extrahiert wurden.

Zudem verfügt sie über Mechanismen, sich selbst zu verbessern, indem die Vorhersage mit dem beobachteten Ereignis abgeglichen wird („auf dem Bild ist tatsächlich eine Katze“). Es geht also, in aller Kürze, um die Wahrscheinlichkeit, mit der eine Aussage als richtig angenommen werden kann. Egal wie die Datenlage auch aussieht, die errechnete Wahrscheinlichkeit ist immer größer als 0 % und immer kleiner als 100 %. Ein Rest an Unsicherheit bleibt immer.

Dabei lassen sich im Wesentlichen zwei Arten von Aussagen unterscheiden: Klassifikation und Vorhersage.



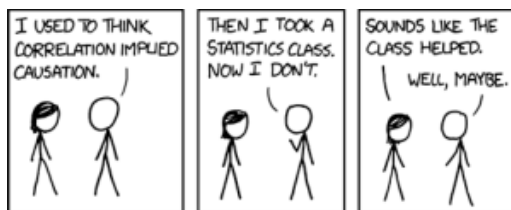
Generische Streudiagramme (scatter plots) für Klassifikation (links) und Vorhersage (rechts), Quelle: Dey, Sankhya. “LINEAR REGRESSION:A Machine Learning Algorithm.” The Startup, July 4, 2020.

Klassifikation heißt, dass einem an sich bedeutungslosen Muster eine Bedeutung zugewiesen wird: Dieses Muster an Pixeln stellt mit Wahrscheinlichkeit X eine Katze dar. Auch hier gilt: Die Wahrscheinlichkeit ist technisch gesehen nie 100 %. Im Fall des obigen Streudiagramms ist die Wahrscheinlichkeit, dass ein Punkt rechts der Linie dunkelgrün (und nicht hellgrün) ist, etwas mehr als 75 %. Wie hoch die Wahrscheinlichkeit X – der sogenannte Schwellenwert (threshold) – sein muss, damit ein System die Antwort als „richtig“ akzeptiert, wird beim Systemdesign festgelegt. Das Gesamtergebnis muss am Ende einfach „gut genug“ sein. Bei der Vorhersage wiederum geht es darum, eine unbekannte

Variable von einer bekannten Variable abzuleiten – etwa die wahrscheinliche Körpergröße auf Grundlage des gemessenen Körpergewichts. Dazu misst man zunächst beide Werte und trägt sie in eine Matrix ein, das Gewicht auf der X-Achse, die Größe auf der Y-Achse. Die Regression reduziert die entstehende Punktwolke auf eine Linie, aus der sich nun zu jedem beobachtbaren Gewicht – jedem X-Wert – die entsprechende Größe – der Y-Wert – „vorhersagen“ lässt. Ob dabei ein Korrelations- oder ein Kausalitätsverhältnis vorliegt, ist mathematisch gesehen ununterscheidbar.

Lernen bedeutet in diesem Rahmen: Jede neue Beobachtung wird zur Grundlage genommen, die Wahrscheinlichkeit einer Aussage neu zu berechnen. Mit jeder neuen Berechnung verschiebt sich die Linie, die das Streudiagramm durchzieht, ein kleines Stück, bis die Passform in den Augen der Systementwickler:innen, beziehungsweise ihrer Auftraggeber:innen, beziehungsweise gemäß den gesetzlichen Vorgaben, wenn es denn solche gibt, gut genug ist.

Korrelationen: Epistemologie der KI



Correlation, Quelle: <https://xkcd.com/552>

Korrelation ist nicht Kausalität – diesen Satz kennt jede und jeder, der oder die je mit Wissenschaft in Berührung gekommen ist. Die moderne Wissenschaft zielt auf kausale Erklärungen ab. A ist die Ursache für B. Durch dieses Wissen ergibt sich die Möglichkeit der Vorhersage, wenn A, dann B.

Sie können uns unterstützen, indem Sie diesen Artikel teilen:

Das epistemische Versprechen der KI ist ein anderes. Man könnte es so formulieren: Aus einer Vielzahl an Korrelationen ergibt sich die bestmögliche Vorhersagewahrscheinlichkeit. Der entscheidende Punkt hier ist die große Menge an beobachteten Korrelationen, die im Einzelnen weder nachvollziehbar noch unmittelbar sinnvoll sein müssen, in der Summe aber eine Aussage ermöglichen sollen. Dabei wird die Erklärung zu Gunsten der Vorhersage

fallengelassen. Statt einer Aussage vom Typ „Wenn A dann B“ erhalten wir solche vom Typ „Es besteht eine x % Wahrscheinlichkeit, dass B eintritt, warum auch immer“. Die Gültigkeit solcher Vorhersagen ist dabei stets zeitlich begrenzt: Sie gilt nur so lange, bis sie auf Basis neuer Daten neu berechnet wird. In der Finanzindustrie, in der viele dieser mathematischen Verfahren ursprünglich entwickelt wurden, kann der zeitliche Horizont auf Sekundenbruchteile zusammenschrumpfen.

Diese Form der Erkenntnisgewinnung hat einige Stärken. Vielleicht die größte: Lernen, in dem hier beschriebenen Sinn, lässt sich automatisieren – Algorithmen können sich selbst anpassen. Damit können Aussagen auf Grundlage unvollständiger Daten gemacht werden, die sich mit dem Eintreffen neuer Daten kontinuierlich verändern, idealerweise verbessern. Auch lassen sich so Aussagen zur Wahrscheinlichkeit nicht wiederholbarer Ereignisse treffen („Der Kurs einer Aktie wird in den nächsten 10 Sekunden steigen.“). Und man kann Aussagen über Systeme machen, deren innere Struktur als Blackbox nicht beobachtbar oder zu komplex und dynamisch für eine Analyse ist.

Diesen Stärken stehen jedoch ebenso bedeutende Schwächen gegenüber. Eben weil das System lernfähig ist, ändern sich Aussagen fortlaufend und sie sind mehrfach kontingent: in Bezug auf die Datenlage, die Ausgangshypothese und den eingestellten Schwellenwert, der eine Aussage als verlässlich einstuft. Und das ist vielleicht die grundlegendste Einschränkung: Mit dieser Methode lässt sich nicht zwischen Wahrscheinlichkeiten und Tatsachen unterscheiden. Wahrheit wird subjektiv: wahr genug, um zu handeln. Wertpapierhändler:innen wollen die Börse nicht analytisch erfassen, sondern zum richtigen Zeitpunkt kaufen oder verkaufen, basierend auf dem eigenen Risikoprofil.

Diese Form der KI halluziniert, im Sinne der fehlenden Unterscheidung zwischen Wahrscheinlichkeit und Faktizität, immer. Nur dass in vielen Fällen der statistische Output den menschlichen Rezipient:innen als richtig erscheint. Nur lässt sich eben nicht vorhersagen, in welchen Fällen das so ist. Deshalb wird die „Intelligenz“ der KI oft als zerklüftet ([jagged](#)) bezeichnet. Spitzen und

Täler liegen oft, in nicht vorhersehbarer Weise, sehr nahe beieinander. Das ist kein Fehler dieser Systeme, sondern eine der Grundeigenschaften der „konnektionistischen KI“ zu der alle heute dominanten LLMs und Diffusionsmodelle gehören.

Analytisch, generativ, agentisch

Aus dieser grundlegenden Funktionsweise – der Berechnung von Wahrscheinlichkeiten auf Basis multipler Korrelationen – wurde in den letzten 15 bis 20 Jahren ein ganzes Spektrum an Anwendungen entwickelt, das sich idealtypisch in drei Grundtypen einteilen lässt: analytische, generative und agentische KI. Diese Unterscheidung ist dabei nicht in erster Linie eine technische, der entscheidende Unterschied liegt vielmehr im Anwendungsfeld. Er zeigt sich am deutlichsten daran, welcher Aussagetypus jeweils produziert wird.

Analytische KI macht Aussagen über Muster, die in den Daten bereits vorhanden sind. Klassische Beispiele sind die optische Zeichenerkennung oder die quantitative Textanalyse. Die Aussage, die hier erzeugt wird, ist eine faktisch-analytische: Sie ist entweder richtig oder falsch. Ein Bild stellt den Buchstaben „S“ dar oder nicht. Analytische Systeme sind dabei in aller Regel auf klar umrissene Aufgabenfelder zugeschnitten. Ein Programm, das darauf trainiert ist, Buchstaben zu erkennen, kann keine Katzen identifizieren.

Generative KI setzt an einem anderen Punkt an. Auch sie beruht auf der Analyse von Mustern in Daten, geht jedoch einen Schritt weiter. Auf der Basis bestehender Muster erzeugt sie neue, ähnliche Muster. Das Neue entsteht durch kontrollierte Momente des Zufalls – weshalb es nur bedingt vorhersehbar und grundsätzlich nicht wiederholbar ist. Damit verschiebt sich auch der Maßstab, an dem man das Ergebnis misst. Es geht nicht mehr um analytische Kriterien, sondern um normativ-ästhetische. Sie lauten entsprechend nicht mehr „richtig oder falsch“, sondern „gelingen oder misslungen“, „schön oder hässlich“.

Die dritte Form ist die agentische KI. Auch sie geht von der Analyse vorhandener Muster aus, nutzt diese aber, um Veränderungen in der externen Welt vorzunehmen. Das Paradebeispiel ist das selbst-

steuernde Auto: Das Modell verändert die Bewegung des Fahrzeugs, beobachtet die Reaktionen der Umwelt und leitet daraus den nächsten Eingriff ab. Das Kriterium, mit dem das Handeln dieser Form der KI beurteilt wird, ist funktional-teleologisch: Wurde das Ziel erreicht?

Für die Frage der generierten Bilder ist vor allem der Übergang vom ersten zum zweiten Typus von Bedeutung. Wo analytische KI noch an einem – wenn auch statistisch erzeugten – Maßstab von richtig und falsch gemessen werden kann, entzieht sich generative KI dieser Bewertung, denn es geht um normativ-ästhetische Urteile. Genau das macht sie, wie sich zeigen wird, für die Frage des öffentlichen Diskurses so voraussetzungsreich.

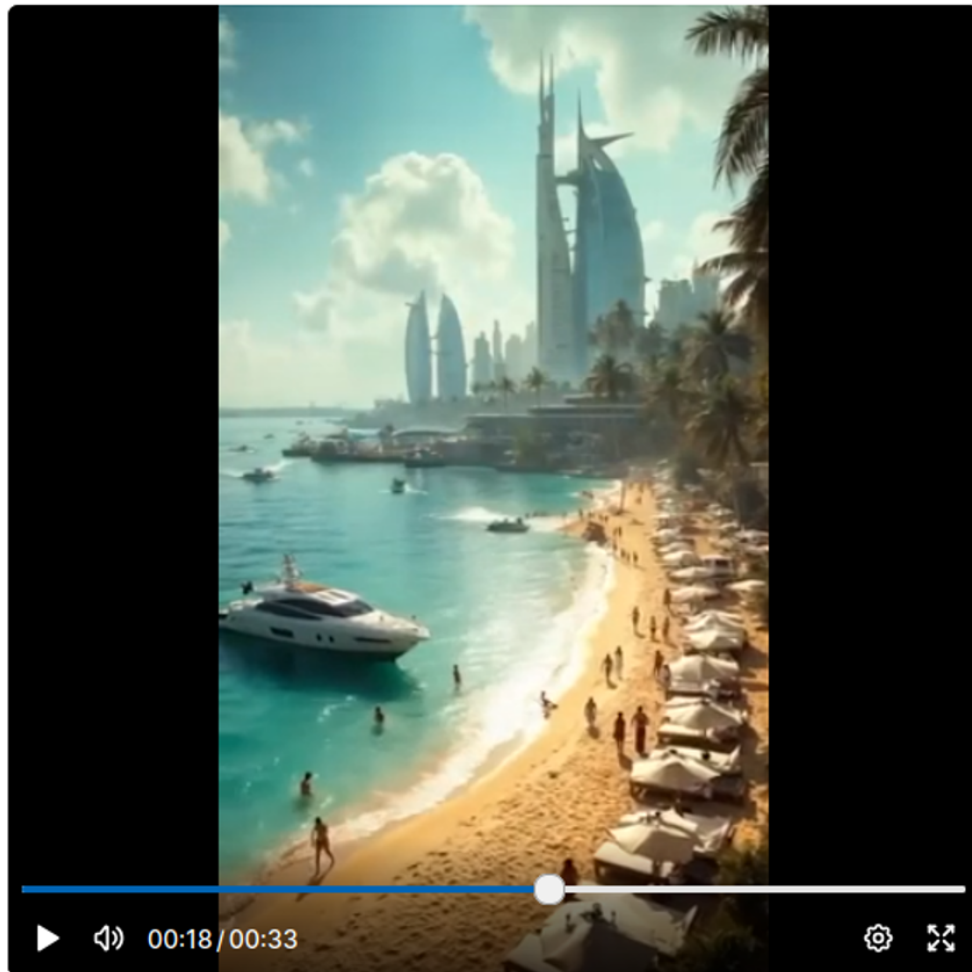
Zeichen und Signale. Generierte Bilder im öffentlichen Diskurs

Viele Kommunikationskanäle werden aktuell überflutet mit generierten Bildern. Manche sind von klassischen Fotos kaum zu unterscheiden, andere klar als solche erkennbar. Doch für beide gilt: Sie sind keine Abbilder einer externen Welt. Was in diesem Zusammenhang als „Realismus“ firmiert, verweist folglich nicht auf ein Verhältnis zur Realität, sondern auf ein Set visueller Repräsentationen und Konventionen, die das Modell reproduziert, ohne sich dabei je auf etwas außerhalb seiner selbst zu beziehen.

Was wir also tatsächlich vor uns haben, wenn wir ein generiertes Bild betrachten, sind statistische Verhältnisse: Die am häufigsten auftretenden Muster eines Datensatzes werden zu einer neuen Konfiguration verdichtet, die es in dieser Form empirisch nicht gibt. Man könnte dieses Verfahren, in Anlehnung an Max Weber, als Bildung eines Idealtypus beschreiben. Ähnlich wie Webers Idealtypus, der ja in seiner begrifflichen Reinheit auf keine konkrete Erscheinung mehr zutrifft, steigern generierte Bilder einseitig bestimmte Regelmäßigkeiten eines Datensatzes und verdichten sie zu einem in sich stimmigen, aber nicht-indexikalischen Bild. Wo Weber jedoch den Idealtypus lediglich als analytisches Hilfsmittel betrachtete, ist das generative Bild das Endprodukt: Es verharrt in der reduktionistischen Klarheit des „Ideals“, das in einem kontrafaktischen Verhältnis zur empirischen Wirklichkeit steht.



Donald J. Trump  
@realDonaldTrump



5.81k ReTruths 19.5k Likes

Feb 26, 2025, 5:51 AM

Quelle: Truth Social, 26.02.2025 (/posts/114068387897265338)

Generative Bilder greifen auf Daten und Muster der Vergangenheit zurück, um daraus potenzielle Zukünfte zu entwerfen: Dinge, die es – noch – nicht gibt, die es aber geben könnte. Die Zielrichtung der Bilder weist damit, bei aller Rückwärtsgewandtheit ihres Ausgangsmaterials, letztlich auf ein utopisches (oder auch dystopisches) Anderswo hin. Die politische Wirkung dieser Bilder im öffentlichen Diskurs ist die Kolonisierung der Zukunft durch die Steuerung der Imagination.

Viele werden sich an das sogenannte „Trump-Gaza“-Video erinnern. Niemand wird es mit der Wirklichkeit verwechseln. Es behauptet auch nicht, die Wirklichkeit darzustellen, sondern entwirft eine visuell konkrete Zukunft des Gazastreifens, in eine Situation hinein, in der niemand weiß, wie diese aussehen könnte. Dass dabei das größte Hindernis dieser Vision, die dafür notwendige

gezielte ethnische Säuberung, bereits als *fait accompli* in der Vergangenheit liegt, ist Teil des Programms. Und dass so viele von uns sich so gut an dieses Video erinnern können, obwohl wir in der Zwischenzeit tausende weitere gesehen haben, spricht für die Macht seiner Wirkung. Als Teil einer politischen Strategie wird die Zukunft hier vorweggenommen, mit dem Ziel, eine spezifische Version denk- und handelbar zu machen. Die Frage, die wir uns stellen sollten, ist keine faktisch-analytische, sondern eine normativ-ästhetische und damit politische: Wollen wir diese Zukunft?

Scrollen statt Lesen

Das von Trump geteilte Video sagt nicht mehr als „Gaza wird Dubai“. Es ist Slop. Aber diese klischeehafte Simplifikation ist genau ihre Stärke, denn diese Bilder (und Videos) sind nicht zum genauen Lesen, sondern zum beiläufigen Scrollen gemacht. Im Kontext sozialer Medien werden generierte, wie indexikalische Bilder in einer solchen Menge und mit einer solchen Geschwindigkeit verbreitet, dass eine inhaltliche Lektüre im eigentlichen Sinn kaum mehr stattfinden kann. Stattdessen werden sie über ihren unmittelbaren affektiven Gehalt rezipiert. Bedeutung entsteht, wenn überhaupt, nicht mehr aus dem einzelnen Bild, sondern über Mustererkennung im Sinne von Wiederholung und Differenz über die Vielzahl der Bilder hinweg, die durch den Feed rauschen.

Diese Mustererkennung findet dabei in zweifacher Hinsicht statt: zum einen auf der Ebene der algorithmischen Filterung, die Produzent:innen im Auge behalten müssen, wenn ihre Bilder überhaupt zirkulieren sollen; zum anderen auf der Ebene der Rezeption durch Nutzer:innen, die sich durch endlose Feeds scrollen. Aus dem Zusammenspiel dieser beiden Ebenen entsteht etwas, das man algorithmische Konformität nennen könnte – die ewig gleichen, doch im Detail immer wieder anders varierten Bilder schöner Menschen, großartiger Reisen, perfekter Wohnungen, oder eben eines Gazastreifens, der aufblüht wie Dubai. Diese Konformität wurde auf Plattformen wie Instagram seit vielen Jahren eingeübt; die generative KI trifft insofern auf bereits ideal vorstrukturierte Bildkulturen.

Damit verändert sich aber die Wirkungsweise der Bilder grundlegend. Semiotisch gesprochen verschiebt sich das Gewicht weg vom Zeichencharakter des Bildes – also von Fragen der Bedeutung – hin zum Signalcharakter. Signale wurden in der Semiotik lange vernachlässigt. Es war vor allem [Félix Guattari](#), der neben bedeutungsgenerierenden, signifikanten Zeichen die Wichtigkeit der affekt- bzw. handlungsgenerierenden, a-signifikanten Signale betonte. Beide Dimensionen sind immer vorhanden, können aber in ihren relativen Gewichten stark variieren. Während erstere über das Bewusstsein vermittelt sind und auf eine Bezeichnung, einen Sinn, zielen, docken letztere unmittelbar an den Körper an, ohne den Umweg über Bewusstsein und Bedeutung zu nehmen. Mit anderen Worten produzieren sie also nicht in erster Linie Sinn, sondern vor allem Affekte, Begehrensweisen, Emotionen, Wahrnehmungen. Auch hier werden zwar Bilder eingesetzt, doch nicht, um Sinn oder Bedeutung hervorzubringen, sondern um eine Handlung, eine Reaktion, ein Verhalten, eine Haltung auszulösen.

Das Ziel ist demnach nicht Indoktrination, sondern Musterverstärkung. Wir befinden uns damit nicht mehr im Feld der klassischen Propaganda, die noch auf Überzeugung, also auf Bedeutung, zielte, sondern im Feld einer Informationstheorie, die bekanntlich ohne Semantik auskommt. Das Resultat lässt sich als eine Art „[Vibe Shift](#)“ beschreiben: Die Zukunft, und mit ihr auch die Gegenwart, fühlt sich plötzlich anders an, ohne dass sich dies auf eine einzelne Aussage zurückführen ließe. Wenn wir von „Slop“ sprechen, meinen wir für gewöhnlich den reduzierten semantischen Gehalt solcher Bilder als Zeichen. Wir übersehen dabei aber leicht, dass sie über eben diese Reduktion, in Verbindung mit ihrer Geschwindigkeit und Wiederholung, gerade an Signalcharakter gewinnen und darüber ihre eigentliche Wirkung entfalten.

Die sozialen Medien sind, was Produktion und Rezeption betrifft, die ideale Umgebung für generierte Inhalte und deren Betonung der Signaldimension. Oder, wie es der Podcaster Vijayendra Mohanty kürzlich formulierte: „[Social media is a fire, AI is kerosene](#)“. Technologisch lässt sich die Produktion von Signalen leicht skalieren, sei es aus kommerziellen, sei es aus propagandistischen Motiven, wobei sich beide in der Praxis oft kaum trennen lassen.

Was wir mit KI-generierten Inhalten also erleben, ist im Kern die Automatisierung der Signalproduktion innerhalb einer algorithmischen Umgebung, die schon vor der generativen KI primär auf Affekt-generierende Signale, nicht auf Bedeutungs-generierende Zeichen ausgelegt war.

Die Gestaltung des Mitschreibens

Auch wenn man annimmt, dass Medien am Denken mitschreiben, und dass KI die Konstitution von Öffentlichkeit verändert, so heißt das nicht, dass hier ein technologischer Determinismus am Werk ist. Technologien werden erst über die Einbettung in gesellschaftliche Prozesse wirkungsmächtig. Die Frage ist also, wie sie eingebettet werden sollen, um den demokratischen öffentlichen Diskurs nicht noch weiter zu schwächen. Klare Antworten darauf kann ich an dieser Stelle nicht anbieten, aber zumindest drei Felder benennen, auf denen ich individuelle und gesellschaftliche Handlungsmöglichkeiten sehe.

Erstes Feld: Anwendung und Grenzen der algorithmischen Epistemologie

Wenn wir KI, wie ich es hier vorgeschlagen habe, als eine spezifische Epistemologie verstehen, dann stellt sich unmittelbar die Frage, für welche Felder und für welche Probleme diese Epistemologie überhaupt angemessen ist. Meine Vermutung ist, dass korrelationsbasierte Verfahren dort am vielversprechendsten sind, wo sie sich auf klar umrissene Anwendungsgebiete mit guter Datengrundlage beschränken, also tendenziell im Bereich der analytischen Anwendungen.

Anstelle der Vision einer KI für alle Zwecke sollten wir uns folglich für die Erhaltung epistemischer Vielfalt einsetzen: für ein Nebeneinander unterschiedlicher, jeweils an ihren Gegenstand angepasster Erkenntnisverfahren, deren Methoden und sozialen Formen der Verhandlung und Vermittlung, anstatt eines einzigen, universell beanspruchten. Es braucht eben nicht nur die Informatik, sondern vielfältige wissenschaftliche und künstlerische Erkenntnisweisen.

Zweites Feld: Die Ebene der Zeichen. Kennzeichnung und Verantwortlichkeit

Ein zweiter Ansatzpunkt liegt auf der Ebene der Lesbarkeit, das heißt der Zeichen. Das bedeutet zunächst, dass man versuchen kann, generierte Bilder als solche durch Kennzeichnung zu identifizieren. Dazu gibt es freiwillige Ansätze der Tech-Firmen, deren praktischer Nutzen aber sehr limitiert ist. Doch selbst wenn, wie im [EU AI Act \(Art 50.2\) gefordert, eine verbindliche Kennzeichnungspflicht für KI-generierte Inhalte umgesetzt werden sollte](#), hilft diese nichts gegenüber jenen, die eine solche Kennzeichnung aktiv umgehen wollen. Und sie trägt nichts zur Frage der Steuerung der Imagination bei, die sich ja gerade nicht auf der Ebene der Zeichen, sondern auf jener der Signale abspielt.

Ein verwandter, aber unterschiedlicher Ansatzpunkt setzt bei der inhaltlichen Verantwortung für den generierten Output an. Im Mai 2026 [entschied das Landgericht München I](#), dass Google als Unternehmen für die Inhalte der KI-Zusammenfassungen von Suchresultaten („AI Overview“) verantwortlich ist und für darin enthaltene Falschaussagen haftet. Wie genau dieses Urteil ausgestaltet wird, ist noch offen. Dieser Weg funktioniert allerdings nur dort, wo sich der Output eines Systems überhaupt am Maßstab von richtig und falsch beurteilen lässt, also im Bereich der analytischen KI. Für generierte Bilder, deren Aussagen primär normativ-ästhetischer Natur sind, greift dieses Modell der Verantwortlichkeit kaum.

Drittes Feld: Die Ebene der Signale. Geschäftsmodell der sozialen Medien

Da die Signalebene der Bilder für ihre gesellschaftliche Wirkung von entscheidender Bedeutung ist, müssen wir an den Mechanismen ihrer Zirkulation ansetzen. Solange soziale Medien ihr Geschäftsmodell auf affektives Engagement ausrichten, wird die Signalebene der Bilder die Ebene der Zeichen strukturell dominieren.

Die Frage ist also, wie ein solcher Eingriff aussehen könnte. Es ist sogar für die EU sehr schwierig, die US-amerikanisch dominierten sozialen Medien regulatorisch von diesem Geschäftsmodell abzubringen oder es auch nur relevant zu beschneiden. Was aber offensteht, ist, uns von diesen Plattformen unabhängiger zu machen. Aktuell wird viel von digitaler Souveränität gesprochen, ohne dass

eine klare Strategie oder wirklicher politischer Wille erkennbar wäre. Dabei wäre es dringend geboten, an anderen Prinzipien orientierte Alternativen zu entwickeln oder die bereits existierenden Ansätze (etwa das dezentrale [Fediverse](#)) konsequenter auszubauen. Ein solches Vorhaben zu forcieren, ist gewiss nicht einfach. Möglicherweise ist der Zeitpunkt dafür aber, Trump und Musk sei Dank, günstiger, als es noch vor kurzem den Anschein hatte.

Welche Rolle generative KI im öffentlichen Diskurs spielen wird, liegt nicht (nur) an der Verfasstheit der KI selbst, sondern vielmehr auch an den gesellschaftlichen Rahmenbedingungen, unter denen sie in soziale Prozesse eingebettet wird. Und diese haben wir selbst in der Hand.

Dieser Text ist eine stark überarbeitete Fassung des Eröffnungsvortrags, den Felix Stalder im Juni 2026 bei der Tagung „Schreibt KI Geschichte?“ an der Gedenk- und Bildungsstätte Haus der Wannseekonferenz gehalten hat.