

Künstliche Intelligenz und Freies Wissen

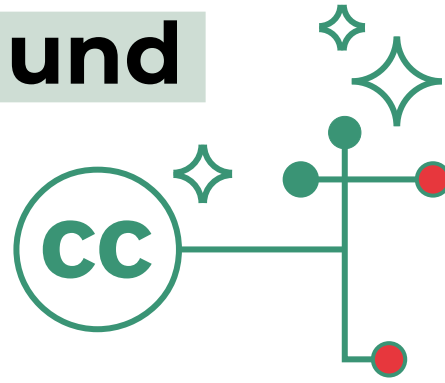


Inhalt

Künstliche Intelligenz und Freies Wissen	4
Warum generative KI das Freie Wissen sowohl fördern als auch schädigen kann	4
Einordnung von KI	4
Nebenwirkungen generativer KI-Systeme und gesellschaftliche Folgen	
der Verengung der Debatte	5
Das Spannungsfeld von generativer KI und Wissen	6
Sprachmodelle erzeugen plausibel anmutende Texte, kein Wissen	6
Das Spannungsfeld von generativer KI und digitalen Gemeingütern	6
Stärkung von Freiem Wissen auch in strukturierter Form	7
Keine echte Offenheit in Aussicht	9
Ist generative KI überhaupt als Gemeingut denkbar?	9
Einseitige Ausbeutung statt Schaffung von Gemeingut	10
Können generative KI-Modelle gemeingutfähig sein?	11
Künstliche neuronale Netzwerke befördern Wissenskonzentration	11
Kriterien für einen angemessenen Umgang mit KI	12
Kriterien für Auswahl und Einsatz	13
konnektionistische KI-Systeme	13
symbolische KI-Systeme	13
Einordnung des KI-Hypes	14
Alternativvorschlag: Verlässliches Wissen als Alleinstellungsmerkmal	14



Künstliche Intelligenz und Freies Wissen



Warum generative KI das Freie Wissen sowohl fördern als auch schädigen kann

Sogenannte Künstliche Intelligenz (KI) ist derzeit ein zentraler Bezugspunkt für viele gesellschaftliche, politische und wirtschaftliche Entwicklungen. Der öffentliche Diskurs spitzt sich dabei hauptsächlich auf zwei Perspektiven zu: blinder Fortschrittsglaube einerseits und dystopische Warnungen andererseits. Beide Sichtweisen verstellen den Blick darauf, dass Menschen – und insbesondere Unternehmen – Technologien gestalten. Sie erschweren damit eine notwendige breite Debatte darüber, welche Gestaltung im Sinne des Gemeinwohls wünschenswert ist. Die häufige Verengung der Debatte auf nur einen Teilbereich der KI – in der Regel die generativen Modelle – verstellt zudem den Blick auf alternative technische Ansätze und deren Potenzial für die Förderung Freien Wissens. Aus genau dieser Perspektive betrachtet Wikimedia Deutschland das Thema.

Zudem begleiten wir neue technische Entwicklungen grundsätzlich offen, setzen uns jedoch zugleich kritisch mit ihnen auseinander. Unser Ziel ist ein Internet, das der gesamten Gesellschaft gleichermaßen dient und bestehende Machtungleichgewichte abbaut. Diese Haltung gründet auf unserem Engagement für freien Zugang zu verlässlichem Wissen, für eine am Gemeinwohl ausgerichtete Digitalpolitik sowie für den Schutz und Ausbau digitaler Gemeingüter (Digital Commons). Wikipedia steht dabei als digitales Gemeingut im Zentrum der Wikimedia-Projekte. Sie alle dienen der Allgemeinheit – nicht nur einzelnen, wenigen –, genau wie ein freies und offenes Netz.

Aus dieser Perspektive stellen sich verschiedene Fragen: Welcher Umgang mit KI ist sinnvoll, um ein verlässliches Informationsökosystem und Freies Wissen zu stärken? Kann KI zu digitalen Gemeingütern beitragen? Welche Rolle kann KI bei der Entwicklung digitaler Gemeingüter spielen? Und welche – bislang zu wenig beachteten – Wege gibt es, um mit etablierten Technologien Wissen besser verknüpfbar und zugänglich zu machen?

Mit diesem Impulspapier möchten wir zeigen: Der aktuell vorherrschende Diskurs rund um KI und ihren Einsatz konzentriert sich nahezu ausschließlich auf generative KI-Modelle. Diese arbeiten probabilistisch, das heißt: Ihre Ausgaben beruhen nicht auf logischen Schlussfolgerungen, sondern auf statistischer Wahrscheinlichkeit. Diese Verengung vermittelt ein irreführendes Verständnis dessen, was Wissen ist und wie es zustande kommt. Bestehende Machtungleichgewichte werden dadurch verstärkt. Für viele Anwendungsfelder, die heute mit generativer KI bearbeitet werden (sollen), existieren deutlich geeignetere Ansätze. Insbesondere die zweite historische Strömung der KI-Forschung kommt hier ins Spiel: die sogenannte symbolische KI. Sie strukturiert gesichertes Wissen in sogenannten Wissensgraphen, die maschinell ausgewertet werden können. Dieser Ansatz ermöglicht verlässliche, deterministische und verlässliche Ausgaben – bei einem Bruchteil des Energiebedarfs generativer Systeme. Zugleich stärkt symbolische KI Freies Wissen und damit digitale Gemeingüter insgesamt.

Einordnung von KI

Der Begriff „Künstliche Intelligenz“ ist selbst innerhalb der Forschung kein einheitlich definiertes Konzept. Seit Beginn des Forschungsfeldes lässt er sich jedoch vereinfacht in zwei Strömungen einordnen: heuristische Modelle, bekannt als konnektionistische KI, sowie logikbasierte Ansätze, die als symbolische KI bezeichnet werden. Erstere arbeiten mit statistischen Methoden und werden mittels maschinellen Lernens iterativ auf Mustererkennung in großen, unstrukturierten Informationsbeständen trainiert. Zweitere folgen einem regelbasierten Ansatz und nutzen Methoden formaler Logik, um beweisbare Schlüsse aus maschinenlesbar strukturierten Wissensbeständen zu ziehen.

In politischen und gesellschaftlichen Debatten steht allerdings der Begriff „KI“ derzeit vor allem für ein Narrativ und wird für eine Vielzahl technischer Verfahren auf unterschiedlichen Ebenen verwendet. In der öffentlichen Debatte wird er meist synonym für generative Modelle zur Text-, Bild- und Videoerzeugung verwendet. Diese Modelle basieren allesamt auf künstlichen neuronalen Netzwerken (KNN), einer Ausprägung des maschinellen Lernens – und damit der konnektionistischen KI.

Klassisches maschinelles Lernen, auch unter Einsatz künstlicher neuronaler Netzwerke, wird bereits seit Langem in verschiedensten Anwendungsbereichen genutzt. Von frühen Spam-Filtern über optische Zeichenerkennung bis hin zu Bilderkennung oder medizinischer Diagnostik geht es dabei in erster Linie um die wahrscheinliche Zuordnung zu einer bestimmten Kategorie – etwa: Handelt es sich um Spam? Welcher Buchstabe ist erkennbar? Ist auf einem Bild ein Hund oder eine Katze zu sehen?

Mit dem sogenannten Deep Learning – also deutlich umfangreicheren KNN in Verbindung mit der Transformer-Modellarchitektur – wurden ab 2017 die Grundlagen geschaffen, um natürliche Sprache zu generieren. Dies bildete einen der Ausgangspunkte für die Entwicklung generativer KI, die heute nicht nur Texte, sondern auch Bilder und Videos erzeugen kann. Sprachmodelle, insbesondere sogenannte Large Language Models (LLMs), machten die Interaktion über natürliche Sprache als Chatbot populär. Bald folgten auch Modelle zum Generieren von Bildern und Videos, die diese inzwischen vertraute Form der Bedienung per Prompt übernahmen. Wenn heute öffentlich über „KI“ gesprochen wird, denkt ein Großteil der Menschen an genau diese generativen Systeme – und an die Interaktionsform über Prompts.

Diese Verengung der Debatte auf lediglich ein Teilgebiet der beiden klassischen Strömungen der Künstlichen Intelligenz hat weitreichende Folgen.

Nebenwirkungen generativer KI-Systeme und gesellschaftliche Folgen der Verengung der Debatte

Da im öffentlichen Diskurs Künstliche Intelligenz derzeit gleichermaßen als abstrakte Projektionsfläche wie auch als Synonym für generative Systeme verwendet wird, konzentriert sich die Debatte vorwiegend auf die Einordnung einer vermeintlich unaufhaltbaren Verbreitung dieser Modelle. Der Vollständigkeit halber möchten wir daher zunächst die unerwünschten Folgen dieser generativen Systeme ansprechen.

Das Training und der Betrieb generativer KI erfordern einen hohen materiellen Ressourcen- und Energieeinsatz. Dadurch ist es nahezu ausschließlich großen Organisationen vorbehalten, eigene Modelle überhaupt zu trainieren. Kleinere Akteure können in der Regel bestenfalls vortrainierte Modelle anpassen – und bleiben somit abhängig von den Anbietern der großen Systeme.

Eine Untersuchung des AI Now Institute schätzt beispielsweise, dass das Training von OpenAIs GPT-4-Modell über 100 Millionen US-Dollar gekostet hat.¹ Hinzu kommen die laufenden Kosten für die Inferenz, also die Nutzung des Modells – diese machen Schätzungen zufolge einen noch größeren Anteil der Infrastrukturkosten großer Cloud-Anbieter aus als das Training selbst.

Auch die Auswirkungen auf die Umwelt sind erheblich: Für den Zeitraum von 2020 bis 2023 wird der durch das Training generativer KI-Modelle verursachte Elektroschrott auf 1,5 bis 5 Millionen Tonnen geschätzt.² Zugleich nehmen KI-Anbieter weiter steigende CO₂-Emissionen in Kauf³ und investieren massiv in zusätzliche Energieerzeugung – etwa durch die Wiederinbetriebnahme stillgelegter Atomkraftwerke,⁴ den Bau von Klein-Kernkraftwerken⁵ und den Ausbau von Gaskraftwerken.⁶

Zudem belastet das massenhafte Training generativer Modelle auch das Web selbst. Die dafür erforderlichen riesigen Mengen an Trainingsdaten – in Form von Texten, Bildern und Videos – werden zunehmend durch aggressive Crawler zusammengetragen. Betreiber*innen von Web-Infrastrukturen berichten zunehmend, dass diese Crawler inzwischen einen signifikanten Teil der Last auf ihren Netzwerken und Servern verursachen. Sie befinden sich nach eigener Aussage in einem dauerhaften Katz-und-Maus-Spiel, um die Belastung in Grenzen zu halten. Dennoch kommt es durch die Crawler immer wieder zu Systemausfällen – mit Folgen auch für reguläre Nutzer*innen der betroffenen Angebote.⁷

So bringt jede Entwicklung und Nutzung generativer KI-Systeme also zwangsläufig unerwünschte Nebenwirkungen mit sich – sei es durch strukturelle Abhängigkeiten von ressourcenstarken Akteuren oder durch den enormen Ressourcenverbrauch, der mit Versuchen verbunden ist, diese Abhängigkeiten durch den Aufbau weiterer Rechenkapazitäten zu verringern. Solche negativen Effekte müssten beim Einsatz generativer Technologien stärker berücksichtigt werden – bislang ist das nur selten der Fall. Sie sind eine weitere Hürde für die Entwicklung generativer KI als digitales Gemeingut und werden selbst dort oft nur am Rande erwähnt. Ein unregulierter, umfassender Einsatz generativer KI droht insbesondere die Umweltauswirkungen – etwa Emissionen, Wasserverbrauch oder Elektroschrott – weiter zu verschärfen. Das steht in offenem Widerspruch zu bestehenden politischen Beschlüssen, insbesondere zu den Zielen der UN für nachhaltige Entwicklung (Sustainable Development Goals).

Das Spannungsfeld von generativer KI und digitalen Gemeingütern

Digitale Gemeingüter wie Wikipedia und andere Wikimedia-Projekte, z. B. die freie Weltkarte OpenStreetMap oder Freie/Open-Source-Software wie der Linux-Kernel, stehen für ein Produktionsmodell, bei dem Güter allen Menschen zur Nutzung, Wiederverwendung, Weiterentwicklung und Pflege zur Verfügung stehen, ohne dass deren rein wirtschaftliche Verwertung im Vordergrund steht.⁸ Alle dürfen diese Projekte und ihre Ergebnisse frei verwenden, weiterverbreiten und auch in veränderter Form nutzen. Sie sind transparent, nachvollziehbar und laden viele Menschen zur aktiven Mitgestaltung ein.

Um sicherzustellen, dass diese Werke als freie Gemeingüter erhalten bleiben und eine einseitige Aneignung oder Ausbeutung vermieden wird, kamen bisher verschiedene Lizenzmodelle auf Grundlage des Urheberrechts zum Einsatz. Dazu gehören etwa die Softwarelizenzen der Freie-Software-Bewegung sowie die Creative-Commons-Lizenzen für nicht-softwarebasierte kreative Werke.

Wir verstehen „Offenheit“ als einen emanzipatorischen Begriff, der darauf abzielt, Machtungleichgewichte abzubauen, Teilhabe zu fördern und der Privatisierung von Wissen entgegenzuwirken. Weder „Offenheit“ noch „Freies Wissen“ oder „Freie Software“ sind dabei ein Selbstzweck – sie sind Mittel und Werkzeuge, um gemeinsame Ressourcen vor einseitiger Ausbeutung zu schützen und gesellschaftliche Teilhabe zu stärken.

Das Spannungsfeld von generativer KI und Wissen

Zugang zu und Teilhabe an Wissen sind zentrale gesellschaftliche Aufgaben. Deshalb steht Freies Wissen im Mittelpunkt der Wikimedia-Projekte. Sie verfolgen das Ziel, gesichertes Wissen zu strukturieren, Querverbindungen herzustellen und möglichst vielen Menschen zu ermöglichen, dieses Wissen frei abzurufen, zu nutzen und weiterzuverwenden. Die Freiwilligen in den Projekten korrigieren Inhalte, wenn Sachverhalte unklar oder verfälscht wiedergegeben werden, und aktualisieren Artikel, wenn neue Erkenntnisse vorliegen oder sich die Einordnung in den Sekundärquellen ändert.

Die aktuelle Entwicklung – speziell im Bereich der Sprachmodelle – befördert hingegen ein problematisches Verständnis von Wissen und trägt vielfach dazu bei, dass Wissensproduktion stärker eingehegt statt allgemein zugänglich gemacht wird. Der öffentliche Diskurs über generative KI, insbesondere in Form von Sprachmodellen, vermittelt ein irreführendes Verständnis dessen, was Wissen ist. Künstliche neuronale Netzwerke (KNN) verstärken dabei die Tendenz zur Wissenskonzentration.

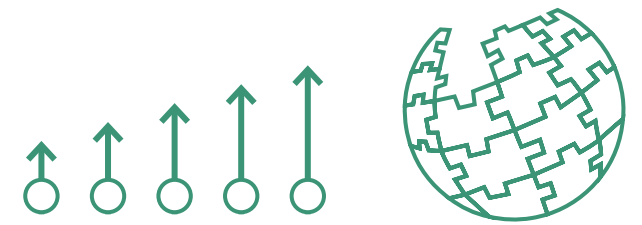
Sprachmodelle erzeugen plausibel anmutende Texte, kein Wissen

Sprachmodelle wie ChatGPT erwecken bei vielen Menschen den irreführenden Eindruck, sie verfügten über menschliche Reflexionsfähigkeit und könnten verlässliches Wissen vermitteln. Auch im Alltag wird häufig formuliert, ein Sprachmodell „wisse“ etwas. Tatsächlich reproduzieren diese Modelle lediglich jene Wortfragmente, die in den Trainingsdaten besonders häufig in der Nähe anderer Wortfragmente vorkommen. Die Ausgaben beruhen auf Wahrscheinlichkeitsberechnungen – ohne dass dabei logische Schlussfolgerungen möglich wären. Das Training bewirkt, dass die generierten Texte für Leser*innen plausibel erscheinen. Es garantiert jedoch nicht, dass sie tatsächlich logisch schlüssig oder inhaltlich korrekt sind.

Der aktuelle Stand der Forschung lässt Grund zum Zweifel, ob große Sprachmodelle oder Chatbots überhaupt geeignete Werkzeuge für die verlässliche Informationssuche und den Abruf von validen Inhalten darstellen. Auch die Fähigkeit zur Überprüfung von Informationen oder zur Einbeziehung zufälliger Funde im Rahmen einer Recherche kann dadurch eingeschränkt sein.⁹ Zwar kann generative KI mit einer gewissen Wahrscheinlichkeit Inhalte ausgeben, die sich bei unabhängiger Prüfung als begründetes, wahres Wissen erweisen. Sie ist jedoch nicht in der Lage, logische Schlüsse zu ziehen – und kann daher auch kein Wissen im klassischen Sinne erzeugen, also keine „begründete, wahre Überzeugung“.

Die Inhalte der Wikipedia kommen auf andere Weise zustande: Ziel ist die Darstellung gesicherten Wissens. Eigene Recherchen der Autor*innen sind in den Artikeln nicht vorgesehen. Stattdessen soll idealerweise der Stand verlässlicher Sekundärliteratur wiedergegeben werden – etwa aus unabhängigen journalistischen Quellen, Fachbüchern oder wissenschaftlichen Veröffentlichungen. Dabei dürfen die urheberrechtlich geschützten Formulierungen dieser Quellen nicht einfach übernommen werden. Vielmehr extrahieren die Autor*innen die zugrundeliegenden (nicht schutzfähigen) Fakten und formulieren daraus eigene Sätze. Kommt es zu unterschiedlichen Einschätzungen über die Bewertung von Quellen, werden diese auf den Diskussionsseiten der jeweiligen Artikel sachlich ausgetragen – mit dem Ziel, eine konsensfähige Formulierung zu finden, die den aktuellen Stand des Wissens möglichst umfassend abbildet.

Die Ausgaben von Sprachmodellen und anderer generativer KI sind also nicht nur plausibel statt korrekt – sie schreiben auch bestehende Muster in die Zukunft fort. Sie reproduzieren die Sprache, mit der sie trainiert wurden – einschließlich enthaltenem Bias, diskriminierenden Mustern und überholten gesellschaftlichen Kategorien. Diese Fortschreibung verengt den Raum für Veränderung erheblich – etwa bei Entscheidungen über Kreditwürdigkeit, über die Zuordnung von Polizeikapazitäten oder über Asyl- und Bewährungsthematiken. Sprachmodelle reproduzieren sprachlich kodierte Vorurteile und Ungleichgewichte in gesellschaftlicher Repräsentation.¹⁰



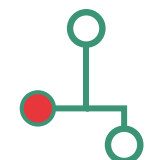
Stärkung von Freiem Wissen auch in strukturierter Form

Alternative Ansätze wie symbolische KI sowie die ihr zugrunde liegenden strukturierten Daten und Wissensgraphen bergen hingegen ein großes, bislang wenig ausgeschöpftes Potenzial: Sie ermöglichen verlässlichere Auskünfte und tragen dazu bei, Wissen umfassender zu erfassen und breiter zugänglich zu machen.

Ein praktisches Beispiel dafür ist Wikidata – eine frei bearbeitbare Wissensdatenbank, in der seit 2012 zehntausende Freiwillige bis heute über 118 Millionen Datenobjekte verfügbar gemacht haben.¹¹ In Wikidata wird Wissen als auswertbarer Wissensgraph strukturiert aufbereitet. Forschungsdatenbanken oder Open-Data-Bestände der öffentlichen Hand und von Unternehmen, die ihr Wissen in ähnlicher Form bereitstellen, können mit der Wissensbasis von Wikidata kombiniert beziehungsweise verlinkt werden. Ein Ziel dieses Ansatzes ist es, einen möglichst großen Teil des menschlichen Wissens in maschinenlesbarer Form zu strukturieren – also als Wissensgraph.

Die maschinelle Auswertung solcher Wissensgraphen erfolgt – im Gegensatz zu generativen Sprachmodellen – nicht induktiv, sondern nach klar definierten logischen Regeln. Das heißt, jede Ausgabe basiert auf eindeutig festgelegten Eingaben, statistische Wahrscheinlichkeiten spielen dabei keine Rolle. Diese Form von Wissensaufbereitung fällt in die Klasse symbolischer KI, die zwar ebenfalls für Automatisierungsprozesse geeignet ist, sich jedoch deutlich vom heuristischen Ansatz der konnektionistischen und insbesondere der generativen KI unterscheidet. Die Wissensbasis von Wikidata und die als Open Data veröffentlichten Wissensgraphen anderer Akteure sind darüber hinaus digitale Gemeingüter und somit Bestandteil des Freien Wissens.

Ein konkretes Beispiel: Eine Großstadt verfügt über zahlreiche faktenbasierte Eigenschaften – etwa die Einwohnerzahl, die Zugehörigkeit zur direkt übergeordneten Gebietskörperschaft, die amtierende Bürgermeisterin und vieles mehr. Sind diese Informationen vollständig in Wikidata erfasst, muss etwa eine geänderte Einwohnerzahl nicht mehr manuell in hunderten Sprachversionen des zugehörigen Wikipedia-Artikels aktualisiert werden. Sie kann stattdessen automatisiert über Wikidata in die Infobox aller relevanter Artikel eingebunden werden.

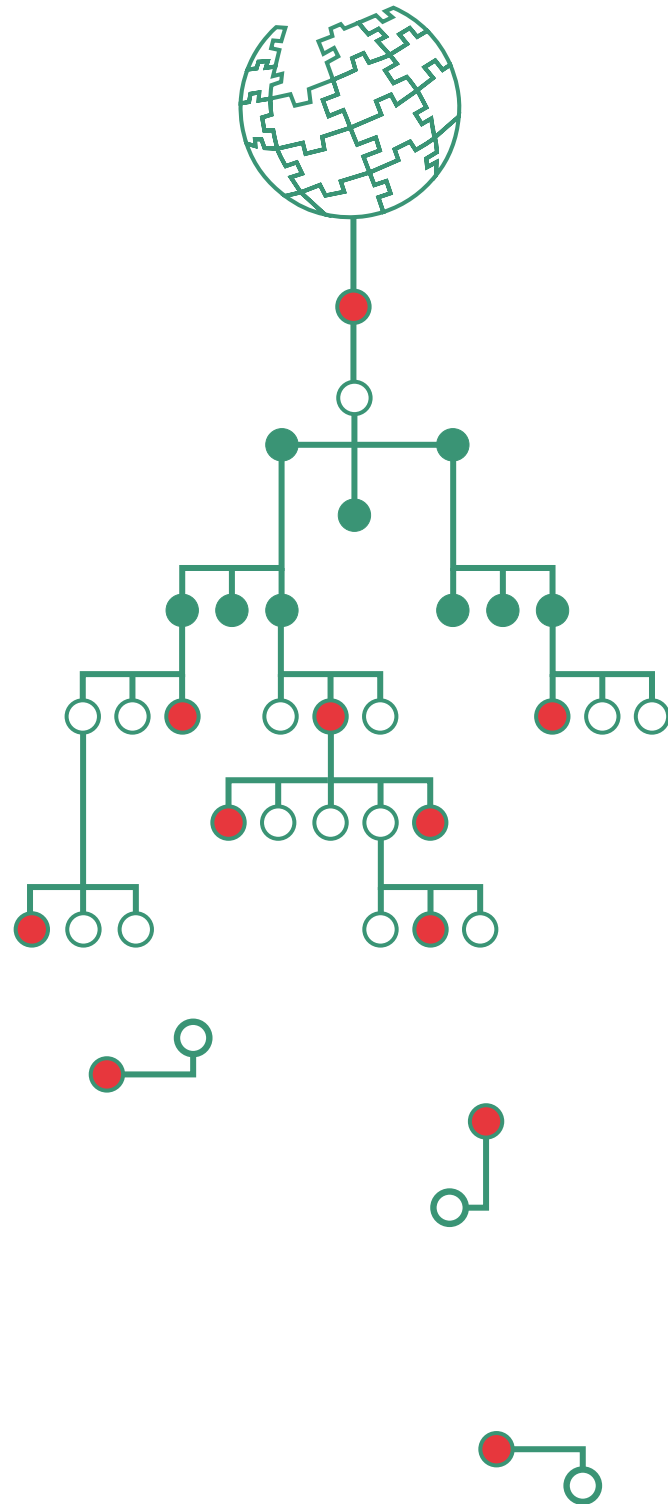


Strukturierte Daten bieten ein erhebliches Potenzial für die Organisation und Abfrage von Wissen – sei es im öffentlichen Sektor, in der wissenschaftlichen Forschung oder in privatwirtschaftlichen Unternehmen. Mit ihnen lassen sich auch komplex scheinende Abfragen leicht umsetzen, beispielsweise: „Welche 20 größten deutschen Städte haben eine weibliche Oberbürgermeisterin?“¹² Eine solche Abfrage ist dank der zugrunde liegenden Datenbasis nachvollziehbar, überprüfbar und belegbar. Unterschiedliche oder fehlerhafte Antworten auf Fragen wie diese – wie sie bei Sprachmodellen häufig vorkommen – sind bei regelbasierten Systemen wie Wikidata ausgeschlossen. Und falls Daten veraltet sind, lassen sie sich in der Regel durch die Nutzenden direkt in der Datenbasis korrigieren – ein wesentlicher Unterschied zu generativen Systemen.

Sprachmodelle hingegen scheitern regelmäßig an solchen Fragestellungen – selbst dann, wenn die relevanten Daten dezentral bereits vorliegen, aber noch nicht an einer Quelle zusammengefasst wurden. Zwar ermöglichen Methoden wie Retrieval Augmented Generation (RAG) die Durchsuchung von Textbeständen, die Ausgabe erfolgt dabei jedoch weiterhin mit statistischer Fehlerwahrscheinlichkeit. Solange aber beispielsweise die Liste der 20 Städte mit Oberbürgermeisterinnen noch nicht explizit als Textquelle existiert, liefert diese Methode kein zuverlässiges Ergebnis.

Wissensgraphen und regelbasierte Systeme bieten hier das Potenzial, sogenannte „unknown knowns“ zu erkennen – also Wissen, über das die Menschheit – oder zumindest eine Organisation – bereits verfügt, das aber bisher nicht so aufbereitet wurde, dass es einfach abfragbar ist. Logikbasierte Systeme arbeiten dabei nicht nur deterministisch und reproduzierbar – sie benötigen für eine Abfrage auch nur einen winzigen Bruchteil der Energie, die ein generatives System für die gleiche Aufgabe verbrauchen würde.

Damit wird deutlich: Die Stärkung des Freien Wissens gelingt vor allem über den Ausbau strukturierter Daten. Generative KI hingegen droht die Vermittlung gesicherten Wissens eher zu erschweren als zu erleichtern. Ob und in welcher Form KI-Modelle außerhalb der Wissensvermittlung einen gesellschaftlichen Beitrag leisten können, sollte auf Basis des Verhältnismäßigkeitsprinzips geprüft werden. Dabei sollten alternative Lösungsansätze ebenso berücksichtigt werden wie die negativen Auswirkungen von KI – auch in solchen Varianten, die als „offen“ bezeichnet werden, es tatsächlich aber nicht sind.



Ist generative KI überhaupt als Gemeingut denkbar?

Trotz der zuvor beschriebenen unerwünschten Nebenwirkungen generativer KI-Systeme lohnt es sich, deren grundsätzliche Eignung als Gemeingut beziehungsweise als digitale Commons zu prüfen.

Die Voraussetzungen für Entwicklung und Betrieb insbesondere großer generativer Modelle deuten stark darauf hin, dass diese auch künftig nur von extrem ressourcenstarken Akteuren hergestellt werden können – oder, in Abhängigkeit von deren Produktionspipelines, bestenfalls angepasst. Das bestehende Web wird dabei als Ressource ausgebeutet; zugleich erfolgt diese Nutzung häufig unter Einsatz prekär beschäftigter Arbeitskräfte im globalen Süden sowie durch hohen Verbrauch natürlicher Ressourcen – etwa Kühlwasser und anderer Rohstoffe für den Betrieb von Rechenzentren.

Gegen die Annahme, dass generative KI-Anwendungen als Gemeingüter entwickelt werden können, sprechen derzeit insbesondere drei Faktoren, die mit der Entwicklung insbesondere der großen generativen Systeme einhergehen:

1. Die einseitige Ausbeutung des Webs, ohne dass dadurch Gemeingüter geschaffen werden, die ebenso gemeinschaftlich entwickelt, verändert und als kollektive Ressourcen begriffen werden können.
2. Die daraus resultierenden, auf generative Systeme fokussierten Abwehrmechanismen, die sich – berechtigterweise – gegen die einseitige Ausbeutung des Ist-Zustands richten, dabei aber Gefahr laufen, bestehende Machtstrukturen eher zu verfestigen als aufzubrechen.
3. Die enorme Ressourcenintensität für die Entwicklung und den Betrieb von großen Modellen in Bezug auf finanzielle Mittel, Personal, Rechenkapazitäten und Infrastruktur – mit weitreichenden Auswirkungen auf Umwelt, Gesellschaft und soziale Gerechtigkeit.

Unter diesen Umständen erscheint es unrealistisch, dass generative KI in einer Weise „offen“ entwickelt werden kann, die dem Gemeinwohl dient, indem sie Machtungleichgewichte ausgleicht und als digitales Gemeingut im Sinne Freien Wissens nutzbar ist.

Keine echte Offenheit in Aussicht

Erhebliche Anstrengungen seitens der KI-Industrie, aber auch von Berater*innen oder Wissenschaftler*innen, zielen derzeit darauf ab, Definitionen von „offener KI“ in öffentlichen Debatten und Gesetzgebungsverfahren zu verankern, die deutlich hinter den seit Jahren etablierten Anforderungen an Offenheit zurückbleiben, wie sie für Freie und Open-Source-Software allgemein anerkannt sind. Diese Bestrebungen haben unter anderem in der europäischen KI-Verordnung ihren Niederschlag gefunden. Dort werden Anreize gesetzt, indem KI-Systeme bevorzugt behandelt werden, „die unter freien und quell-offenen Lizenzen veröffentlicht werden“. Solche Systeme müssen geringere regulatorische Anforderungen erfüllen als andere.¹³ Was diese Regelung und die damit eingeführten Anreize jedoch konkret bedeuten, bleibt in der Verordnung unklar.

Die Bezeichnung solcher Systeme als „offen“ verkörpert dabei – mehr oder minder ausgeprägt – sogenanntes Openwashing: **Hersteller bezeichnen ihre Modelle oder Systeme als „Open-Source-AI“, ohne dabei dem zentralen Anliegen Freier Software zu folgen – nämlich Machtasymmetrien abzubauen, indem Software unabhängig reproduzierbar gemacht wird.** Derzeit vorgeschlagene Definitionen von „Open Source AI“¹⁴ verändern das bisherige von Freier und Open-Source-Software gewohnte Verständnis von Offenheit. Überträgt man diese beiden Definitionen auf Software, könnte auch Software, könnten auch Programme als „Open Source“ gelten, deren Quellcode nur teilweise offengelegt ist, die nur als vorgefertigtes, kompilierbares Paket genutzt werden können und die sich nur in engen Grenzen anpassen lassen. Dieses Offenheitsverständnis verändert nicht nur die bisherigen Definitionen von Freier und Open-Source-Software, sondern auch die mit dieser Offenheit verbundenen gesellschaftlichen Ziele.^{15 16}

Zur Bewertung, ob ein KI-Modell als „offen“ gelten kann, unterscheiden Forschende insgesamt 14 Kriterien – darunter den Zugang zum Quelltext des Modells, zur Trainingspipeline und zu den Trainingsdaten. Weitere Kriterien für echte, unabhängige Reproduzierbarkeit eines Modells sind etwa die Veränderbarkeit der Parametergewichtungen und der Methoden zur Feinanpassung sowie die Dokumentation des Trainingsprozesses.¹⁷ In der Praxis beschränkt sich die postulierte „Offenheit“ jedoch meist auf die Möglichkeit, ein fertiges, vortrainiertes Modell und möglicherweise die Parametergewichtun-

gen herunterzuladen, das Modell zu verwenden und es neu zu parametrisieren.¹⁸ Selbst diese Möglichkeiten zur Nutzung sind oft an Bedingungen geknüpft – etwa an Registrierungspflichten oder Einschränkungen bei der Weiterverwendung.¹⁹

Besonders widersprüchlich erscheinen diese Beschränkungen vor dem Hintergrund des juristischen Rahmens: Die Hersteller generativer Systeme gehen davon aus, dass 1) die Nutzung der zum Training notwendigen Trainingsdaten mit Verweis auf die Schrankenbestimmungen für Text- und Datamining urheberrechtlich unbedenklich ist, 2) auch die Ausgaben eines generativen Modells nicht urheberrechtlich geschützt sind, aber 3) die trainierten Modelle selbst mit Verweis u.a. auf das von ihnen vorgenommene signifikante Investment beim Training in Verbindung mit dem Leistungsschutzrecht für Datenbankhersteller sehr wohl geschützt sind – denn nur so könnten sie die von ihnen vorgegebenen Beschränkungen bei der Wiederverwendung ihrer Modelle überhaupt durchsetzbar machen.

Nach der OSI-Definition müssen Trainingsdaten nicht zwingend vollständig offengelegt werden. Sie erlaubt beispielsweise ausdrücklich die Verwendung nicht-öffentlicher Daten, sofern sie beschrieben werden. Auch das „Model Openness Framework“ erfordert lediglich in der am weitesten gefassten Klassifikation („Open Science“), dass alle für eine unabhängige Reproduktion eines Modells notwendigen Trainingsdaten tatsächlich verfügbar sind. Für unabhängige Dritte – etwa im Rahmen von IT-Sicherheitsaudits – ist es daher praktisch unmöglich, diese Modelle zu reproduzieren. Entgegen der seit vielen Jahren bewährten Definition für Freie und Open-Source-Software bedeutet „Offenheit“ in diesem Kontext also keineswegs, dass die Reproduktion eines Softwaresystems prinzipiell allen zugänglich sein sollte, um einseitige Machtkonzentrationen zu vermeiden. Der Vollständigkeit halber muss angemerkt werden, dass selbst bei vollständig quelloffenen Modellen eine praktische Reproduzierbarkeit für eine breite Öffentlichkeit angesichts des dafür notwendigen Ressourcenaufwands und der benötigten Datenvolumina ohnehin kaum umsetzbar. Es ist daher fraglich, ob eine breite Skalierung dieses Prinzips unter dem Vorzeichen der Offenheit gesellschaftlich überhaupt wünschenswert wäre.

Einseitige Ausbeutung statt Schaffung von Gemeingut

Ein idealisiertes Bild für die Erzeugung von Gemeingütern ist das einer Allmende, die von Gleichberechtigten geschaffen wird. Das bedeutet: Alle Beteiligten entscheiden eigenständig und gleichberechtigt, was sie zum Gemeingut beitragen können – und dieses Gemeingut soll letztlich allen zugutekommen.

Das Ziel ist also nicht allein eine Nutzenoptimierung des fertigen Produkts, sondern dabei auch sicherzustellen, dass der Herstellungsprozess auf Augenhöhe erfolgt – also ohne einseitige Bevorzugung bestimmter Akteursgruppen zulasten anderer.²⁰ Eine solche gemeinschaftliche Produktion ist bei der Entwicklung generativer Modelle bislang nicht erkennbar. Die Aneignung kollektiver Ressourcen erfolgt zudem teils gegen den ausdrücklichen Willen von Urheber*innen.

Dass in der aktuellen KI-Entwicklung eher Wert abgeschöpft als geschaffen wird, zeigt sich besonders deutlich an zwei Bereichen: der Debatte um mögliche Änderungen des Urheberrechts und der prekären Clickwork im globalen Süden. Wertabschöpfung bedeutet, dass Unternehmen Gewinne aus Ressourcen oder Tätigkeiten generieren, die sie selbst nicht bereitstellen.²¹ Intensiv diskutiert wird derzeit, ob und in welchem Umfang die für das Training herangezogenen Daten durch Urheber- oder verwandte Schutzrechte geschützt sind oder geschützt werden sollten.²² Um einem potenziellen Schutz zuzukommen, schließen KI-Entwickler*innen gezielt Nutzungsverträge mit besonders laut- und durchsetzungsstarken Akteuren wie etwa dem Online-Forum Reddit²³ oder dem Medienkonzern Axel Springer²⁴. Weniger sichtbar, aber nicht minder relevant ist die Tatsache, dass für das Annotieren von Daten oder das Anpassen von KI-Outputs Menschen in Kenia und anderen Ländern des globalen Südens unter äußerst schlechten Arbeitsbedingungen tätig sind.²⁵

Angesichts der schieren Größe von Trainingssätzen großer KI-Modellen erscheint es kaum vorstellbar, dass Entwickler*innen zugleich für eine legale, konsensuale Nutzung der Trainingsdaten sorgen und zugleich hinaus Mechanismen für Teilhabe und Mitgestaltung bereitstellen könnten.

Können generative KI-Modelle gemeingutfähig sein?

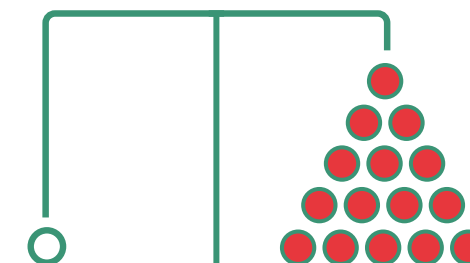
Die derzeitige Entwicklung und der Betrieb großer generativer KI-Modelle erfüllen die Anforderungen an digitale Gemeingüter nicht. Umso wichtiger ist es, die damit verbundenen negativen Wirkmechanismen klar zu benennen, ihnen entgegenzuwirken und einen möglichen Nutzen generativer KI-Modelle sorgfältig gegen ihre negativen Effekte abzuwägen.

Es ist allerdings nicht absehbar, dass generative KI in absehbarer Zeit überwiegend gesellschaftlichen Zielen dienen kann. Daher ist Vorsicht geboten gegenüber der Vorstellung, „weniger schädliche“, „etwas offenere“ oder „vergleichsweise legal trainierte“ KI-Systeme seien bereits wünschenswerte Alternativen, die gesellschaftlich gefördert werden sollten.

Diese Vorsicht gilt insbesondere für Vorschläge, große KI-Modelle als staatlich betriebene Infrastruktur der Daseinsvorsorge aufzubauen. Zwar würden solche Modelle Alternativen zu privatwirtschaftlichen Angeboten darstellen. Die grundsätzlichen Probleme generativer Systeme, von mangelnder Offenheit bis zu den Seiteneffekten der Produktionsmethoden, blieben dann jedoch weiterhin bestehen.

Künstliche neuronale Netzwerke befördern Wissenskonzentration

Einzelne Anwendungen – insbesondere im Zusammenhang mit künstlichen neuronalen Netzwerken (KNN) – können der Forschung dabei helfen, neues Wissen zu generieren. Dieses Wissen wird jedoch häufig durch Patente oder andere Schutzmechanismen abgesichert, um es kommerziell verwerten zu können. Ein Beispiel hierfür ist AlphaFold von Google: Das KNN ermöglichte einen Durchbruch in der Vorhersage von Proteinstrukturen – einem zentralen Feld der medizinischen Forschung. Trainiert wurde es mit umfangreichen öffentlichen Daten.²⁶ Doch Google hat die Offenheit des Codes und der Gewichtungen unter Verweis auf kommerzielle Interessen zunehmend eingeschränkt²⁷ und erst nach öffentlichem Druck wieder partiell geöffnet – allerdings weiterhin mit Einschränkungen, insbesondere für nicht-akademische oder kommerzielle Nutzung.²⁸



Kriterien für einen angemessenen Umgang mit KI

Wir plädieren dafür, im Diskurs über Künstliche Intelligenz eine präzise Begriffsnutzung zu etablieren. Das bedeutet: Es sollte klar benannt werden, welche Art von KI-System jeweils gemeint ist und welche Schlussfolgerungen sich daraus ergeben. Für welche Ziele und Zwecke eignet sich ein System grundsätzlich – und für welche nicht? Welche Machtverhältnisse werden durch den Einsatz eines KI-Systems verstärkt, welche potenziell abgebaut? Diese Fragen sind auch entscheidend für die Bewertung, welche Form von „Offenheit“ einem bestimmten System oder Konzept tatsächlich zugeschrieben werden kann. „KI“ sollte Sinne dieser präzisen Begriffsverwendung nicht lediglich als Sammelbegriff für das Teilfeld des maschinellen Lernens oder gar nur für generative Systeme verwendet werden.

Wir sind davon überzeugt, dass Technologie – und damit natürlich auch KI – in erster Linie dem Gemeinwohl dienen soll. Die Auswahl geeigneter Technologien sollte daher an verschiedenen Kriterien orientiert sein, die das Gemeinwohl stärken.³¹ Dazu gehören unter anderem:

1.

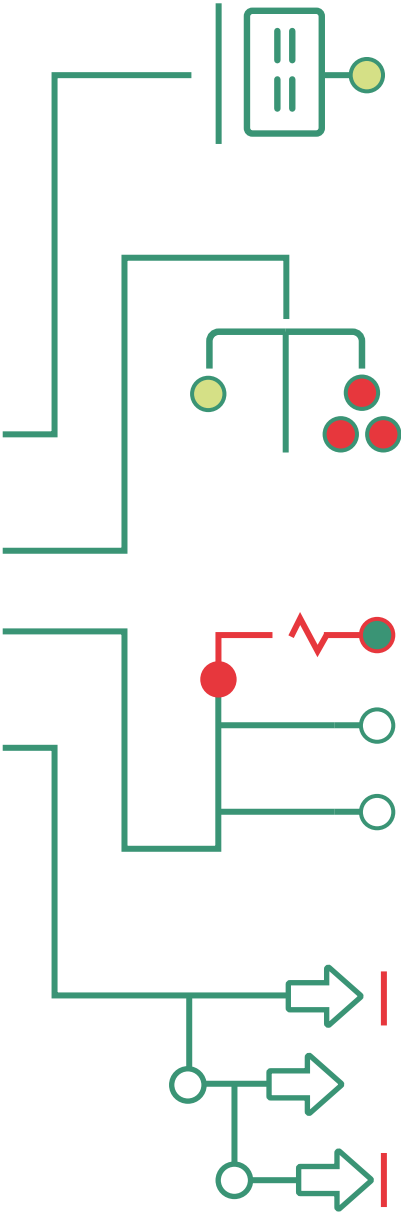
Wie transparent und nachvollziehbar sind die Entstehung und potenziellen Auswirkungen eines Systems?
2.

Trägt es zur Verringerung von Ungleichheit bei – oder verstärkt es bestehende bzw. schafft sogar neue?
3.

Wie zugänglich ist die Technologie? Wer kann sie nutzen, verändern, zu ihr beitragen? Dient sie allen Teilen der Gesellschaft gleichermaßen?
4.

Fördert oder behindert sie künftige Entwicklungen, Innovationen und gesellschaftlichen Wandel?

Im Folgenden zeigen wir auf, wie ein angemessener, zielorientierter Umgang mit unterschiedlichen Formen von Künstlicher Intelligenz aussehen kann. Dazu gehört insbesondere die Unterscheidung, welche Klassen von KI-Systemen für welche Aufgabenstellungen überhaupt berücksichtigt werden könnten. Es zeigt sich deutlich: Für das Bereitstellen von und den Zugang zu Wissen gibt es vielfach geeignetere Technologien als generative KI – allen voran die strukturierte Aufbereitung in Form von Wissensgraphen, die sich automatisiert und verlässlich auswerten lassen.



Kriterien für Auswahl und Einsatz

konnektionistische KI-Systeme	symbolische KI-Systeme
→ Bilderkennung, Zeichenerkennung, in Verbindung mit Sprachmodellen z. B. auch Handschriftenerkennung in Archiven o.ä. → Ausformulierung von Texten auf Basis von Stichworten – mit anschließender Prüfung → Audiotranskription oder -übersetzung (heute oft lokal auf einem Laptop möglich) Generell: Aufgaben mit „verrauschten“ Eingaben und offenen Problemstellungen (Open-World-Annahme), bei denen stochastische Fehler entweder leicht erkennbar oder prinzipiell tolerierbar sind.	→ Situationen, in denen Ausgaben nachvollziehbar und logisch geschlossen begründet sein müssen → Formalisierte Verfahren mit hohem Korrektheitsanspruch, z. B. Verwaltungsverfahren (Stichwort: Gleichheitssatz) → Verlässlicher und nachvollziehbarer Wissensabruf Generell: Information Retrieval, Beweisbarkeit, logische Schlussfolgerungen, Closed-World-Annahme, stochastische Fehler sind nicht akzeptabel.

Um allgemein den Sinn vom Einsatz von Automatisierungssystemen gleich welcher Art zu beurteilen, schlagen wir vier Prüfschritte vor. Diese orientieren sich am bekannten Prinzip der Verhältnismäßigkeitsprüfung:³²

1.

Ist der Zweck der Technologie **legitim**? In einem gesellschaftlichen Kontext sollte neben der rein rechtlichen Beurteilung auch eine Rolle spielen, inwieweit die Zwecke mit gesellschaftlichen Normen und Werten vereinbar ist. Die Beantwortung dieser Frage fällt insbesondere dann schwer, wenn – wie bei vielen generativen KI-Modellen – gerade die unklare Zweckbestimmung als Vorteil gilt.
2.

Ist das untersuchte Mittel **geeignet**, um dieses Ziel zumindest prinzipiell zu erreichen? Es geht hierbei allein um die Frage, ob die Technologie zur Zielerreichung beitragen kann.
3.

Ist das untersuchte Mittel **erforderlich**, um das Ziel zu erreichen? Hier stellt sich die Frage, welche anderen und möglicherweise geeigneteren Mittel es gibt, die beispielsweise verlässlichere Ergebnisse liefern, weniger Ressourcen bedürfen oder gar beides. Vielfach werden derzeit konnektionistische KI-Systeme – vor allem generative Modelle – für Zwecke eingeführt, für die es bereits - in mehrerlei Hinsicht - besser geeignete Lösungsansätze ohne den Einsatz generativer KI gibt.
4.

Ist der Einsatz des untersuchten Mittels **angemessen**? Eine Abwägung der erhofften Vorteile der Technologie mit den in Kauf zu nehmenden weiteren Effekten wie Ressourcenverbrauch, Konflikten mit bestehenden Beschlusslagen z.B. zu den Sustainable Development Goals, sowie Konflikte mit Selbstbestimmungszielen bei Digitalisierungsvorhaben zeigt auf, ob der erwartete Erfolg die damit verbundenen Seiteneffekte rechtfertigen könnte.
- Eine Analyse fehlgeschlagener „KI“-Projekte legt unter anderem den Schluss nahe, dass vielfach das angestrebte Ziel vorab gar nicht klar definiert wurde, oder dass es bei den Projekten weniger um das Ziel an sich ging, sondern um den Selbstzweck, ein „KI-System“ einzusetzen.³³ Das hier vorgeschlagene Prüfschema kann helfen, von Beginn an für Klarheit über das angestrebte Ziel zu sorgen – und in einem zweiten Schritt zu analysieren, ob bereits erprobte Technologien existieren, mit denen sich dieses Ziel besser erreichen lässt.

Einordnung des KI-Hypes

Die aktuellen Entwicklungen rund um generative KI vollziehen sich vor dem Hintergrund massiver Kapitalinvestitionen: Große Tech-Konzerne investieren derzeit mehrere hundert Milliarden US-Dollar in technische Infrastruktur wie Rechenzentren – mit der Erwartung, diese Ausgaben künftig gewinnbringend zu verwerten.³⁴ Zugleich sehen sich politische Entscheidungsträger*innen auf verschiedenen Ebenen in einem vermeintlichen Wettbewerb – sei es zwischen Weltregionen oder innerhalb nationaler Ressorts – um die schnellere und umfassendere Nutzung generativer KI. Die Sorge, in diesem Wettlauf „abgehängt“ zu werden, ist weit verbreitet.³⁵

Angesichts der großen Bedeutung, die vor allem generativen KI-Systemen beigemessen wird, ist es für Politik und Gesellschaft essenziell zu erkennen: **Entwicklung und Einsatz von KI sind gestaltbar.** Dies gerät allerdings oft aus dem Blick, nicht zuletzt, weil Medien, Unternehmen oder Fachleute suggerieren, dass sich insbesondere generative KI zwangsläufig durchsetzen werde.³⁶ Dabei wird verkannt: KI ist keine autonome Kraft mit eigenen Interessen oder anderen menschlichen Eigenschaften. Wenn KI etwa als „halluzinierend“ oder „arbeitend“ beschrieben wird, entsteht der Eindruck eines handelnden Akteurs – was verdeckt, dass es sich hierbei um von Menschen bzw. Unternehmen entwickelte und gesteuerte Systeme handelt.

Die Frage lautet daher nicht, wie sich eine angeblich unausweichliche, vollständige Durchdringung der Gesellschaft durch generative KI-Systeme lediglich noch in Feinheiten gestalten lässt. Vielmehr gilt es, auch Vorschläge zu sogenannten „etwas offeneren“ oder „gemeinwohlorientierteren“ Formen generativer KI-Systeme kritisch einzuordnen. Der Begriff „Public AI“ wird in diesem Kontext häufig als Gegenmodell zur Abhängigkeit von wenigen Konzernen präsentiert – verbunden mit der Hoffnung, dadurch, meist nur diffus umrissene, Innovationspotenzfelder zu erschließen.³⁷ Doch angesichts des Spannungsverhältnisses zwischen konnektionistischen KI-Systemen auf der einen und der Definition von Wissen als gesichert und nachvollziehbar auf der anderen Seite bleibt fraglich, ob solche Ansätze tatsächlich zur Förderung von gesellschaftlichen Zielen wie Offenheit und Freiem Wissen beitragen können.³⁸

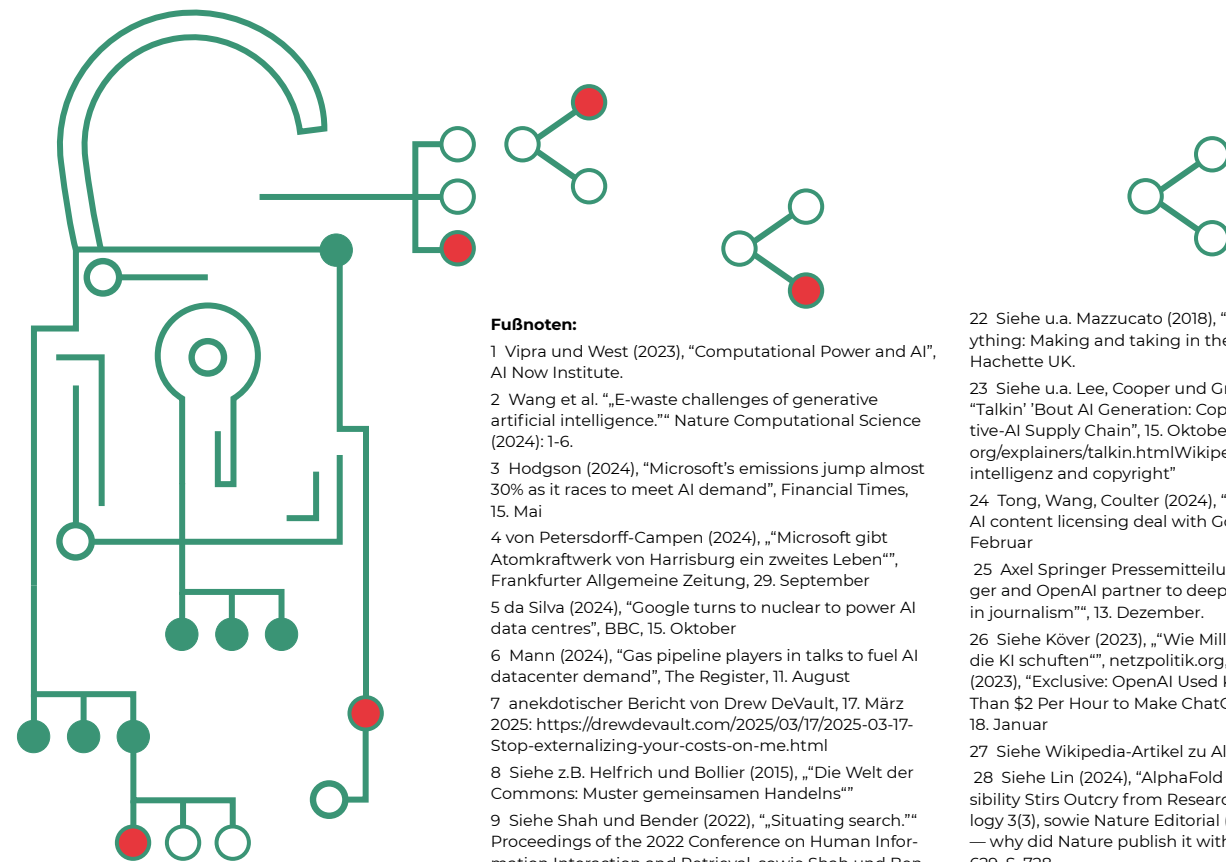
Alternativvorschlag: Verlässliches Wissen als Alleinstellungsmerkmal

Wir stellen daher die Annahme in Frage, dass es politisch und strategisch überhaupt wünschenswert ist, sich vorrangig auf den „Wettbewerb“ rund um die Entwicklung und Anwendung generativer KI-Systeme einzulassen – mit dem Ziel, in diesem Feld zu konkurrieren oder Alternativen zu schaffen. Gerade mit Blick auf die verlässliche, überprüfbare Ausgabe von gesichertem Wissen gerät durch diesen Wettbewerb das enorme Potenzial symbolischer KI-Systeme aus dem Blick.

Diese Systeme brillieren nicht nur durch ihren um Größenordnungen geringeren Ressourcen- und Energiebedarf. Die für ihren Einsatz nötigen Wissensgraphen und strukturierten Datenbanken führen außerdem dazu, dass die Veröffentlichung von Freiem Wissen und offenen Daten nicht mehr als zusätzlicher Aufwand, sondern als nahezu automatisches Nebenprodukt entsteht. Den größten Nutzen erzielen dabei zunächst die Institutionen selbst, die ihre Informationsbestände in dieser Form aufbereiten: Sie profitieren von einem besseren internen Zugriff, verlässlicheren Auswertungen, der Verknüpfbarkeit bislang disparat in verschiedenen Office-Dokumenten abgelegten Informationen – und der Möglichkeit, viele bislang aufwändig händisch vorzunehmender Prozesse zu automatisieren.

Die Veröffentlichung jener Teile der aufbereiteten Wissensbestände, die sich als Open Data eignen, wird so zu einem einfachen, automatisierten Schritt im Anschluss an die interne Aufbereitung. Da sich die unterschiedlichen, auf diese Weise veröffentlichten Wissensgraphen aufeinander beziehen lassen und gemeinsam sowie dezentral ausgewertet werden können, entsteht ein Gesamtbestand an Freiem Wissen, der mehr ist als die bloße Summe seiner Teile: Jeder einzelne Beitrag zu diesem Gesamtbestand erhöht zugleich den Wert der bereits vorhandenen und nutzbaren Informationen.

Statt sich also in das vermeintliche Wettrennen um die besten generativen Systeme einzureihen – oder zumindest ergänzend dazu – könnte es ein starkes Alleinstellungsmerkmal sein, sich als Vorreiter für gesichertes, verlässlich auswertbares Wissen zu positionieren.



Fußnoten:

- Vipra und West (2023), "Computational Power and AI", AI Now Institute.
- Wang et al. "E-waste challenges of generative artificial intelligence." Nature Computational Science (2024): 1-6.
- Hodgson (2024), "Microsoft's emissions jump almost 30% as it races to meet AI demand", Financial Times, 15. Mai
- von Petersdorff-Campen (2024), "Microsoft gibt Atomkraftwerk von Harrisburg ein zweites Leben", Frankfurter Allgemeine Zeitung, 29. September
- da Silva (2024), "Google turns to nuclear to power AI data centres", BBC, 15. Oktober
- Mann (2024), "Gas pipeline players in talks to fuel AI datacenter demand", The Register, 11. August
- anekdotischer Bericht von Drew DeVault, 17. März 2025: <https://drewdevault.com/2025/03/17/2025-03-17-Stop-externalizing-your-costs-on-me.html>
- Siehe z.B. Helfrich und Bollier (2015), "Die Welt der Commons: Muster gemeinsamen Handelns"
- Siehe Shah und Bender (2022), "Situating search," Proceedings of the 2022 Conference on Human Information Interaction and Retrieval, sowie Shah und Bender (2024), "Envisioning information access systems: What makes for good tools and a healthy Web?" ACM Transactions on the Web 18.3, S. 1-24.
- Siehe zuletzt Shojaaee, Parshin, et al. (2025) "The illusion of thinking: Understanding the strengths and limitations of reasoning models via the lens of problem complexity." arXiv preprint arXiv:2506.06941.
- Siehe Koteck, Dockum, Sun (2023), "Gender bias and stereotypes in large language models", Proceedings of the ACM collective intelligence conference, S. 12-24 sowie Hofmann, Kalluri, Jurafsky, et al. (2024), "AI generates covertly racist decisions about people based on their dialect", Nature 633, S. 147-154.
- Stand: Juni 2025, die aktuelle Zahl kann auf der Startseite des Projekts https://www.wikidata.org/wiki/Wikidata:Main_Page eingesehen werden
- Beispielabfrage: <https://w.wiki/5VAq>
- KI-Verordnung der EU, Art. 2 Abs. 12 und Art. 53
- z.B. "The Open Source AI Definition" (OSAID) der Open Source Initiative oder das "Model Openness Framework" (MOF) der Linux Foundation – 1.0 R-C1
- Zur Definition Freier Software siehe z.B. "What is Free Software?" des GNU-Projekts, zur Definition von Open-Source-Software die "Open Source Definition" der Open Source Initiative. Beide zu finden im Wikipedia-Artikel zu "Freie Software"
- U.a. Jürgen Geuter (2024), "Does Open Source AI really exist?", 16. Oktober.
- Liesenfeld und Dingemanse (2024), "Rethinking open source generative AI: open washing and the EU AI Act", Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency (FAccT '24), S. 1774-1787.
- Widder, Gray und Whittaker (2023), "Open (for business): Big tech, concentrated power, and the political economy of open AI"
- Tarkowski (2023), "The Mirage of Open-Source AI: Analyzing Meta's Llama 2 release strategy", Open Future. Der Download der Modelle über Metas Webseite erfordert die Übermittlung von Kontaktdaten an Meta für den Zugang zu den Llama-Modellen, das gleiche gilt für den Download über die Modell-Plattform HuggingFace.
- Auch die Mission der Wikimedia-Projekte ist nicht lediglich die Verfügbarkeit von Wissen, sondern dass Menschen dabei "ermächtigt und eingebunden" sind, siehe <https://wikimediafoundation.org/about/mission/> ("to empower and engage people around the world").

22 Siehe u.a. Mazzucato (2018), "The Value of Everything: Making and taking in the global economy". Hachette UK.

23 Siehe u.a. Lee, Cooper und Grimmelmann (2023), "Talkin' 'Bout AI Generation: Copyright and the Generative-AI Supply Chain", 15. Oktober <https://blog.genlaw.org/explainers/talkin.html> Wikipedia-Artikel zu "Artificial intelligence and copyright"

24 Tong, Wang, Coulter (2024), "Exclusive: Reddit in AI content licensing deal with Google", Reuters, 22. Februar

25 Axel Springer Pressemitteilung (2023), "Axel Springer und OpenAI partner to deepen beneficial use of AI in journalism", 13. Dezember.

26 Siehe Köver (2023), "Wie Millionen Menschen für die KI schuften", netzpolitik.org, 17. März und Perrigo (2023), "Exclusive: OpenAI Used Kenyan Workers on Less Than \$2 Per Hour to Make ChatGPT Less Toxic", Time, 18. Januar

27 Siehe Wikipedia-Artikel zu AlphaFold

28 Siehe Lin (2024), "AlphaFold 3 Angst: Limited Accessibility Stir Outcry from Researchers", GEN Biotechnology 3(3), sowie Nature Editorial (2024), "AlphaFold3 — why did Nature publish it without its code?", Nature 629, S. 728

29 Nuñez (2024), "Google DeepMind open-sources AlphaFold 3, ushering in a new era for drug discovery and molecular biology", VentureBeat, 11. Oktober

30 Rikap, Lundvall (2020), "Big tech, knowledge predation and the implications for development", Innovation and Development, 12(3), S. 389-416 sowie Rikap (2024), "Intellectual monopolies as a new pattern of innovation and technological regime", Industrial and Corporate Change, 33(5), S. 1037-1062.

31 Siehe Themenseite zu Wissensgerechtigkeit: <https://www.wikimedia.de/themen/wissensgerechtigkeit/>

32 Siehe das Politikpapier von Wikimedia zu Gemeinwohl in der Digitalpolitik (2023)

33 Die Verhältnismäßigkeitsprüfung wird normalerweise bei Eingriffen in Freiheitsrechte verwendet. Das Schema eignet sich aber unserer Meinung nach auch gut für eine Prüfung in diesem Fall.

34 Ryseff, De Bruhl, Newberry (2024), "The Root Causes of Failure for Artificial Intelligence Projects and How They Can Succeed: Avoiding the Anti-Patterns of AI", RAND Research Report.

35 Allein im Jahr 2024 beliaufen sich die Investitionen von vier großen Tech-Unternehmen (Alphabet, Amazon, Microsoft und Meta) in physische Infrastruktur wie Rechenzentren und Chips auf ca. 200 Milliarden USD, siehe The Economist (2024), "Big tech's capex splurge may be irrationally exuberant", 16. Mai.

36 U.a. Kaltheuner, Saari, Kak, West (2024), "Reorienting European AI and Innovation Policy", AI Now Institute, [<https://think.ging.com/articles/ai-monthly-competition-intensifies-and-europe-falls-behind/>]

37 U.a. Charr (2024), "The Inevitable AI: Art Of Growth With Generative Intelligence", Stardom Books, und Sample (2024), "An AI Fukushima is inevitable: scientists discuss technology's immense potential and dangers", The Guardian, 22. November.

38 Siehe Marda, Sun, Surman (2024), "Public AI: Making AI work for everyone, by everyone", Mozilla Foundation, und Jackson, Cavello, Devine, Garcia, Klein, Krasodomski, Tan, Tursman (2024), "Public AI: Infrastructure for the Common Good", Public AI Network.

39 Wikimedia Deutschland spricht sich dafür aus, KI in der Bildung konsequent offen auszugestalten, was die unerwünschten Nebeneffekte von einem Einsatz von KI in der Bildung erheblich verringert, siehe https://www.wikimedia.de/wp-content/uploads/2024/04/Handlungsempfehlungen_Offene_KI_fuer_alle.pdf

Über Wikimedia Deutschland

Wikimedia Deutschland ist ein gemeinnütziger Verein mit über 111.000 Mitgliedern und 180 Beschäftigten, der sich für die Förderung von frei verfügbarem Wissen im digitalen Raum einsetzt. Als größte Ländervertretung der internationalen Wikimedia-Bewegung fördert der Verein die ehrenamtlichen Communitys der Wikipedia und weiterer Wikimedia-Projekte in Deutschland. Wikimedia Deutschland entwickelt und pflegt freie Software und die freie Datenbank Wikidata. Der Verein engagiert sich im digital- und bildungspolitischen Bereich für Rahmenbedingungen, die den freien Zugang zu Wissen und Daten möglich machen. Zudem kooperieren wir mit Kulturinstitutionen, um mehr kulturelles Erbe frei zugänglich zu machen.

Der Text entstand in Zusammenarbeit mit Aline Blankertz.



Bei Fragen oder Gesprächsbedarf zu den Inhalten kontaktieren Sie gerne Stefan Kaufmann
Referent Politik und öffentlicher Sektor
stefan.kaufmann@wikimedia.de

Impressum

Wikimedia Deutschland e. V.
Tempelhofer Ufer 23/24
10963 Berlin

Telefon: +49 (0) 30 577 11 620

Geschäftsführende Vorständin

Franziska Heine
Eingetragen im Vereinsregister des Amtsgerichts
Charlottenburg,
VR 23855

Redaktion

Franziska Kelch

Inhaltlich verantwortlich

Lilli Iliev

Gestaltung

Rasmus Giesel, MOR Design,
www.mor-design.de

Die Texte und das Layout dieser Broschüre werden unter den Bedingungen der »Creative Commons Attribution«-Lizenz CC BY-SA in der Version 4.0 veröffentlicht.
<https://creativecommons.org/licenses/by-sa/4.0/deed.de>

Alle Links in dieser Publikation wurden letztmalig am 19. Juni 2025 abgerufen

