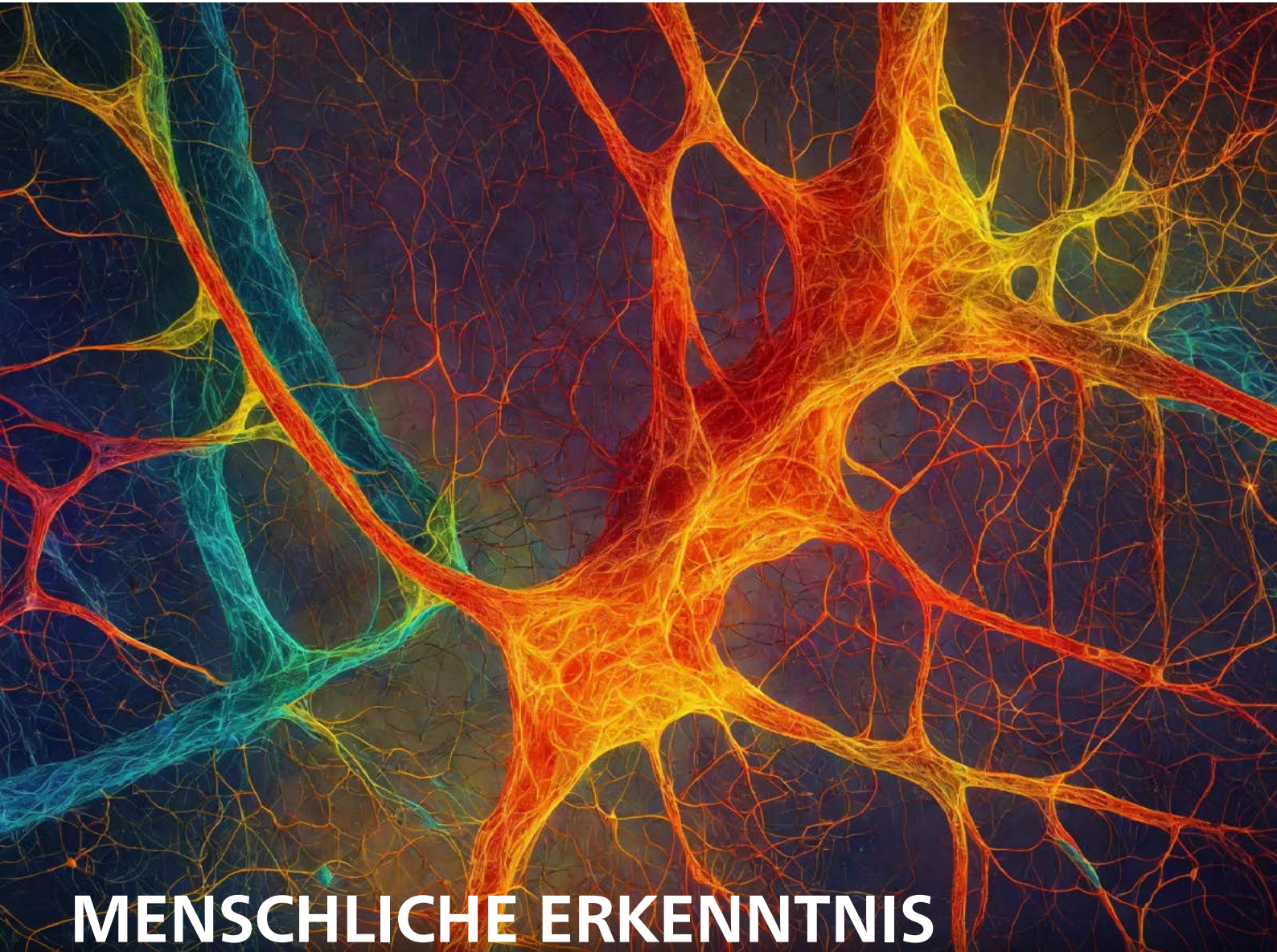




Forum

für **Universität und Gesellschaft**

Universität Bern



MENSCHLICHE ERKENNTNIS IN DER ÄRA KÜNSTLICHER INTELLIGENZ

Veranstaltungsreihe
Winter 2024/25

u^b

b
**UNIVERSITÄT
BERN**

Editorial

Liebe Leser:innen

Seit der Einführung von Chat GPT im September 2022 hat Künstliche Intelligenz ein breites Publikum erreicht. Innerhalb kurzer Zeit begannen Millionen von Anwender:innen, die Kapazitäten dieser umfangreichen Antwortmaschine zu erforschen, wodurch sie gleichzeitig als freiwillige Tester zur Optimierung des Systems beitrugen.

Generative KI-Modelle sind darauf ausgelegt, menschliche Interaktionen nachzuahmen. Ihre Präsenz prägt unsere Ära: Alle grossen Technologieunternehmen bieten mittlerweile eigene Lösungen zur Text- und Bildgenerierung an. Bei den Nutzer:innen führen diese Systeme zu einer Verstärkung bestehender Trends: Technologiebegeisterung trifft auf dystopische Ängste.

Durch KI ist es möglich, riesige Datensätze weit über menschliche Kapazitäten hinaus zu analysieren. Die genauen Funktionsweisen bleiben jedoch oft unklar. Wie kann man KI-Modelle also dazu bringen, das Gewünschte zu tun? Wer ist verantwortlich bei Schäden? Sind unsere Daten sicher? Lassen sich KI-Systeme regulieren?

Das Forum für Universität und Gesellschaft bietet Orientierung. Es veranschaulicht anhand von Beispielen, wie KI funktioniert und in welcher Weise sie unseren Alltag bereits beeinflusst. Dabei werden sowohl die sich bietenden Chancen und Möglichkeiten als auch die ethischen und rechtlichen Herausforderungen beleuchtet, die jetzt angesprochen werden müssen. Ziel ist es, ein tieferes Verständnis für KI zu entwickeln und für die damit verbundenen Problematiken zu sensibilisieren.

Alle Referate wurden aufgezeichnet und sind hier mit kurzen Zusammenfassungen repräsentiert. Das Play-Symbol verweist auf die entsprechenden Inhalte.

Wir wünschen Ihnen eine gute Lektüre.

Forum für Universität und Gesellschaft



Inhalt

Menschliche Erkenntnis in der Ära Künstlicher Intelligenz: <i>Marcus Moser</i>	4
«Den Algorithmen ist die Demokratie scheisseegal»: <i>Eduard Kaeser</i>	6
Was ist KI und wie funktioniert sie?	10
Was ist Künstliche Intelligenz und denkt sie wirklich? <i>Claus Beisbart</i>	
Das Gehirn als Vorbild: <i>Mihai A. Petrovici</i>	
Wie Maschinen lernen: <i>Sigve Haug</i>	
Bilder bearbeiten, Bilder manipulieren: <i>Petra Müller</i>	
Die dunkle Seite der Macht?	14
Algorithmen selektionieren politische Informationen: <i>Silke Adam</i>	
Eine neue Dimension der Propaganda: <i>Jennifer Victoria Scurrrell</i>	
Nutzen und Grenzen von KI: <i>Ruth Fulterer</i>	
Neuronale Abläufe mit Lu.i verstehen	18
Künstliche Intelligenz für natürliche Gesundheit?	20
Krankheiten überwinden dank KI? <i>Thomas Lemmin</i>	
Potential und Grenzen von Chatbots in der Psychotherapie: <i>Thomas Berger</i>	
Bessere Diagnostik dank KI: <i>Piotr Radojewski</i>	
Roboter zur Unterstützung älterer Menschen: <i>Alexander Seifert</i>	
Ethische Implikationen beim KI-Einsatz: <i>Stephen Milford</i>	
Die Zukunft? Bildung und Regulation	26
KI und Kreativität – Eine Frage der Anerkennung: <i>Hannes Bajohr</i>	
KI und die Universität – Vom Wissen zum Können: <i>Fritz Sager</i>	
KI – Macht ohne Kontrolle? <i>Estelle Pannatier</i>	
KI-Regulierung – Die Schweiz setzt auf Zurückhaltung: <i>Sabine Brenner</i>	
KI und Wissen – Zwischen Präzision und Täuschung: <i>Tobias Hodel</i>	
Künstliche Intelligenz in Anwendung	32
Häufig verwendete Begriffe kurz erklärt	33
Referent:innen / Moderation / Projektgruppe	35
Impressum	36



Das Play-Symbol verweist auf die entsprechenden Video-Inhalte.
Sämtliche Aufnahmen und Unterlagen zur Veranstaltung sind
auf unserer Website abrufbar:
www.forum.unibe.ch/kuenstlicheintelligenz

Menschliche Erkenntnis in der Ära Künstlicher Intelligenz



Künstliche Intelligenz fordert natürliche Intelligenz ganz gewaltig heraus. Sie fasziniert – und macht Angst. Der *Homo sapiens* kommt unter Druck, der Mensch kann seine Bedeutung aber als *Homo sentiens*, als fühlender Mensch, unterstreichen.

Von Marcus Moser

Der Begriff «Künstliche Intelligenz» ist schon recht alt. Er wurde Mitte der 1950er Jahre vom US-Informatiker John McCarthy kreiert. In jenen Jahren wurden erste Konzepte zur Simulation menschlicher Intelligenz formuliert. Aber erfahrbar war sie trotz gelegentlichen Sensationsmeldungen über menschliche Niederlagen (Schach oder das Strategiespiel GO) für die breite Bevölkerung nicht. Das hat sich Ende November 2022 drastisch verändert. Damals wurde ChatGPT von der in San Francisco ansässigen Firma OpenAI als Web-App veröffentlicht.

Das Sprachmodell schlug ein wie die sprichwörtliche Bombe, denn mit einem Mal konnten sich viele aktiv mit Künstlicher Intelligenz befassen und herausfinden, was diese zu leisten im Stande ist. Und das Ganze sprachbasiert – und damit in einem Medium, das wir alle kennen. Bereits zwei Monate nach dem Start hatte die App über 100 Millionen Nutzerinnen und Nutzer – eine unglaubliche Zahl.

Das – meine Damen und Herren – das ist dem *Homo sapiens* ganz schön in die Knochen gefahren. Wie war das nun mit hohen kognitiven Fähigkeiten, mit Sprache, Kultur und all den technologischen Fertigkeiten, die dem *Homo sapiens* seit je zugeschrieben wurden?

Neue Kränkung für den *Homo sapiens*

Der Psychoanalytiker Sigmund Freud sprach 1917 von drei menschlichen Kränkungen, die – so seine These – das Selbstverständnis des Menschen und seine Stellung in der Welt fundamental in Frage stellen. Ausgelöst wurden diese Kränkungen allesamt durch wissenschaftliche Entdeckungen, die zu weitreichenden Paradigmenwechseln führten.

Sie erinnern sich: Die erste Kränkung war die *Kopernikanische*. Nicht die Erde steht im Mittelpunkt des Universums. Sie ist lediglich ein Planet, der um die Sonne kreist. Was für ein Schlag fürs menschliche Selbstbewusstsein.

Und weiter geht's mit der *zweiten Kränkung – der Darwinistischen*. Der Mensch ist nicht Produkt einer göttlichen Schöpfung sondern das Ergebnis eines langen, natürlichen Evolutionsprozesses. Der Mensch ist von den Tieren nicht grundsätzlich zu unterscheiden, sondern wie sie ein Produkt der gleichen biologischen Prozesse. Sie wissen, dass diese Auffassung auch heute noch bestritten wird.

Damit kommen wir zur *dritten, zur Freudschen Kränkung* – die eben von Sigmund Freud selber stammte: Das Unbewusste regiert das menschliche Verhalten. Der Mensch ist nicht vollständig Herr und Frau seiner eigenen Gedanken und Handlungen. Er ist mitgesteuert von unbewussten Trieben und Prozessen.

Freud sah in diesen drei wissenschaftlichen Entdeckungen tiefe Verletzungen des menschlichen Egos, da sie das Selbstverständnis des Menschen als *Mittelpunkt der Welt*, als *einzigartiges Geschöpf* und als *rationales Wesen* infrage stellten. Jede dieser Kränkungen führte zu einem neuen Verständnis der Welt und des Menschen, was seine besondere Rolle und Bedeutung relativierte.

Inzwischen werden Sie wohl erahnen, worauf ich hinaus will. *Ich möchte hier die Annahme formulieren, dass das Auftauchen und die Funktionsweise von Künstlicher Intelligenz als vierte Kränkung des Menschen betrachtet werden kann.*

Warum? Die Begründung liegt in der Infragestellung der *menschlichen Einzigartigkeit* und *Überlegenheit* in Bereichen, die lange als Domäne des menschlichen Verstandes galten:

Bisher war eben der Mensch derjenige, der für *komplexes Denken, Kreativität, Problemlösung und Entscheidungsfindung* verantwortlich war.

Mit der Entwicklung von KI-Systemen, die in bestimmten Aufgabenbereichen schneller, präziser und oft besser als Menschen sind, wird die menschliche intellektuelle Überlegenheit infrage gestellt. KI kann Muster erkennen, aus Daten lernen und Entscheidungen treffen – ähnlich, wie es ein Mensch tun würde. Der entscheidende Unterschied ist, dass sie das auf der Basis immenser Datenmengen und mit enormer Geschwindigkeit tun kann.

Und wir können uns dem nicht entziehen – im alten Bild: Der Geist ist aus der Flasche, die Künstliche Intelligenz ist da. Und zwar vielenorts: In der Sprachverarbeitung und Übersetzung, in der Bild- und Videoverarbeitung, bei Empfehlungssystemen, im Gesundheitswesen, bei autonomen Fahrzeugen, im Finanzwesen, in der Industrie und Produktion, aber auch in der Bildung, der Landwirtschaft sowie im kreativen Bereich.

Das eröffnet enorme Chancen – birgt aber auch immense Risiken. Stellen Sie sich ein Blatt mit einer Trennungslinie vor – die klassische Plus-/Minus-Liste. Zu jedem der genannten Anwendungsbereiche könnten Sie mehrfach Chancen und korrespondierende Risiken auflisten. So kann die Automatisierung von Routineaufgaben zur Entlastung führen, aber auch zum Jobverlust. Die Liste liesse sich beliebig fortsetzen.

Wir werden hier von beidem sprechen und von beidem sprechen müssen. Damit erfüllt das Forum für Universität und Gesellschaft gewissermassen den ersten Schritt zu einem verantwortungsvollen Umgang mit KI: Aufklärung und Schulung.

Gefordert werden weitherum Regulierungen und ethische Leitlinien, um den Einsatz von KI verantwortungsvoll zu gestalten.

Das Thema ist ganz schön anspruchsvoll, etwas unheimlich und kann Ängste auslösen. Um so wichtiger ist es, möglichst viel davon zu verstehen. Verstehen gibt uns das Gefühl, die Situation besser einschätzen und beeinflussen zu können. Wenn wir die Ursachen und möglichen Folgen eines Phänomens kennen, wird es weniger bedrohlich.

Homo sentiens – der fühlende Mensch

Womit wir wieder beim *Homo sapiens* wären. Max Tegmark, Professor für Physik am MIT und KI-Experte, schlägt uns ein RE-BRANDING vor. *Statt als Homo sapiens empfiehlt er uns die Selbstbezeichnung als Homo sentiens – als fühlender Mensch.*

Damit hebt er die Bedeutung von Emotionen, Empathie, Bewusstsein und ethischem Denken hervor. Während KI den *Homo sapiens*, den «wissenden» oder rational denkenden Menschen, in vielen Bereichen ergänzt oder sogar übertrifft, bleibt der Mensch als emotionales und bewusstes Wesen einzigartig.

Diese Verschiebung könnte den Menschen dazu anregen, sich stärker auf seine emotionalen und moralischen Fähigkeiten zu konzentrieren und seine Rolle als kreativer und empathischer Akteur zu betonen. Schauen wir uns die Weltlage an, wäre dies wahrlich wünschenswert.

«Den Algorithmen ist die Demokratie scheissegal»

Der Autor diskutiert in seinem Essay die Beziehung zwischen Algorithmen, Demokratie und Gesellschaft. Seine zentrale These lautet: Algorithmen sind indifferent gegenüber demokratischen Werten, können jedoch aufgrund ihrer Struktur und Nutzung politische und soziale Macht entfalten.

Von Eduard Kaeser

Der Titel meines Beitrages mag Sie etwas nassforsch anmuten. Aber er stellt für mich so etwas wie eine Reverenz vor einem Philosophen dar: dem Amerikaner John Haugeland, den ich für einen der tiefsten Denker über Künstliche Intelligenz halte. Er prägte den Satz: Dem Computer ist die Welt scheissegal («they don't give a damn»). In diesem Sinn also meine ich: Den Algorithmen ist die Demokratie scheissegal. Darüber möchte ich mich im Folgenden etwas auslassen.

Zunächst einmal gilt die Umkehrung: Autokratien sind Algorithmen nicht scheissegal. Sie sehen in der KI-Technologie vielmehr ein Wundermittel zur Festigung ihrer Macht. Es überrascht kaum, dass die Kommunistische Partei Chinas die KI-Forschung als «einen neuen Schwerpunkt des internationalen Wettbewerbs» bezeichnet. Die Vermutung ist naheliegend, dass die Priorität nicht im wissenschaftlichen Erkenntnisfortschritt liegt, sondern im Wettrennen der Überwachung. Hier hat sich China schon als führend etabliert. Unsummen werden in die KI-Forschung investiert. China exportiert zudem seine Überwachungstechnologien in rund 80 Länder. Vorzugsweise in solche mit schwachen demokratischen Strukturen. Über Tiktok hat das kommunistische Regime auch Zugriff auf Daten von westlichen Nutzer:innen.

Technik und Politik

Ich will mich nicht in China-Bashing ergehen. Mich beschäftigt hier etwas anderes. Und wiederum zitiere ich einen amerikanischen Philosophen, Langdon Winner. Er stellte 1980, also vor dem Siegeszug der digitalen Technologie, in einem Artikel die Frage «Do artifacts have politics?»: «Machen Artefakte Politik?». Sie klang damals höchst seltsam, ja dissonant in der Technikdiskussion, galt doch die Technik nach vorherrschender Meinung als «neutral» – als Mittel zum Zweck. Dass das technische

Mittel selbst politisch wirkmächtig sein und das Verhalten des Menschen steuern kann, war ein ziemlich unorthodoxer Gedanke. Denn man argumentierte vorzugsweise «anthropozentrisch»: Der Mensch bestimmt, das Gerät gehorcht. Wie nun aber Winner schrieb: «In unserer Zeit sind die Menschen oft bereit, ihre Lebensgewohnheiten zu ändern, nur um sich technischen Innovationen anzupassen; indes würden sie ähnliche Änderungen ablehnen, begründete man sie politisch».

Genau dies lässt sich von der neuen digitalen Technologie sagen. Sie setzte sich nicht demokratisch durch, sondern stillschweigend und schleichend als Verhaltensänderung durch Technik. Denken wir nur an die Smartphones. Wir haben unser Verhalten spielend ihrer «Politik» angepasst. Und indem wir das tun, konditioniert uns die Technologie der Sozialen Medien immer mehr.

«Die neue digitale Technologie setzte sich nicht demokratisch durch, sondern stillschweigend und schleichend als Verhaltensänderung durch Technik.»

Algorithmen als Akteure

Wie also können Algorithmen eine politische Wirkmacht ausüben? Es geht, wie gesagt, nicht primär um die Frage, was politische Akteure mit Algorithmen anstellen, sondern wie die Maschinen selbst als politische «Akteure» operieren. Ohne sie vermenschlichen zu wollen, gestehen wir ihnen heute eine Eigeninitiative zu, wie wir



dies ja bei Menschen ganz selbstverständlich tun. Eine lose Metaphorik spricht schon seit langem von Automaten, die «entscheiden», «planen», «wahrnehmen», «auswählen», «voraussehen». Umgangsformen greifen Platz, in denen Maschinen die Rolle von Partnern, Akteuren, Quasi-Personen, Usurpatoren, womöglich Feinden übernehmen. Die Maschine beginnt zu lernen. Sie entwickelt sich auf eine kognitive Stufe zu, auf der man ihr menschenähnliche Intelligenz zuschreibt – bald einmal vielleicht Bewusstsein, Intentionalität, Sensibilität, ja, Persönlichkeit. Und damit beginnt sie sich in unseren gesellschaftlichen Körper zu integrieren. Sie wird zivil. Einigen Robotern hat man bereits Bürgerrechte zugesprochen. Das Tool wird zum Mitbürger. Wir bewegen uns also auf eine Homo-Robo-Gesellschaft zu.

Genauer gesehen, stecken wir schon mittendrin. Am deutlichsten zeigt sich dies in der Debatte um Kontrolle und Transparenz der neuen KI-Systeme. Dass diese Tools immer auch die Signatur ihrer Designer tragen, ist schon seit längerem bekannt. Man spricht vom KI-Bias, den Vorurteilen und Voreingenommenheiten, die sich in den Programmen niederschlagen und zu Fehlurteilen führen können. Das gilt besonders für die Deep-Learning-Systeme. Ihre Aufgabe ist, Muster in den Trainingsdaten zu erkennen, und Muster sind primär dann erkennbar, wenn sie oft genug in Erscheinung treten. Stereotype sind repetierte Muster. Ohne Ironie lässt sich deshalb sagen: Generative KI generiert vorzugsweise Stereotype, und in diese sind fast definitionsgemäss Vorurteile und Voreingenommenheiten eingegossen. Die einschlägigen Technologiekreise sind sich dieses Problems durchaus bewusst. Und sie suchen nach Tools des Bias-Testens, IBM etwa mit Watson OpenScale, Google mit What-If. Microsoft gibt eine Liste kritischer Algorithmusstudien heraus.

Es existiert ein neues spezielles Computerforschungsfeld namens «fairness, accountability and transparency in machine learning (FATML)».

Dass man «Gerechtigkeit», «Verantwortlichkeit», «Transparenz» von Maschinen und ihren Algorithmen einfordert, lässt sich durchaus als ein ethisches Erwachen der Computerwissenschaftler deuten, die sich allmählich der sozialen und politischen Folgen ihrer KI-Systeme bewusst werden. Es existiert eine Vielzahl solcher Probleme. Ihnen entspricht die wichtige Aufgabe, demokratische Institutionen zu schaffen, die solchem Missbrauch entgegenzutreten. Im Mai 2023 verabschiedete die EU einen regulatorischen Rahmen für den Einsatz von Algorithmen. Wir brauchen aber mehr, nämlich eine Perspektive, die Szenarien des gesellschaftlichen Computereinsatzes analysiert. Hier gilt es in erster Linie nicht die neuen technischen Applikationen – also die klassifizierenden und die generativen KI-Tools – ins Visier zu nehmen, sondern die generelle Tendenz, soziale und politische Probleme ingenieural zu definieren – was sich letztlich genau als Kernproblem erweist: Kann man fehlerhafte Algorithmen stets mit noch mehr und besseren Algorithmen flicken? Gerechtigkeit, Verantwortlichkeit, Transparenz sind ja nicht primär technische Begriffe, sondern politische, lassen sich also nicht optimieren wie eine Maschinenfunktion. Ein derartiger Technozentrismus nimmt leicht technokratische Züge an. Und in der Technokratie lauert immer die Autokratie.

Opazität der Maschine

Wenn schon von Transparenz die Rede ist – man spricht heute von der Opazität der Maschine, also von ihrer Undurchschaubarkeit. Hier muss man mindestens drei Facetten unterscheiden, die Demokratie-Risiken bergen: die technische, die institutionelle und die individuelle. Neuronale

Netzwerke bestehen heute schon aus Milliarden von Parametern, die einen numerischen Input in einen numerischen Output verwandeln. Das heisst, das System entwickelt womöglich Prozeduren, die nur es selbst kennt. Der Designer hat begrenzten – wenn überhaupt – Einblick in das, was sich im Innern abspielt. Mit zunehmender Schichttiefe wird das Netz selbständiger: eine Black Box. Der Technikhistoriker George Dyson hat sogar ein Gesetz formuliert: Ein System, das man vollständig versteht, ist nicht komplex genug, um intelligentes Verhalten zu manifestieren; und ein System, das intelligentes Verhalten manifestiert, ist zu komplex, um vollständig verstanden zu werden.

Die institutionelle Form der Undurchsichtigkeit gefährdet die Demokratie am offensichtlichsten. Viele marktgängige KI-Systeme sind geschützt durch Firmengeheimnis. Wir können zwar die Produkte nutzen, zum Beispiel die GPT-Modelle. Sie funktionieren lokal auf den eigenen Computern. Aber man braucht dazu die Verbindung zu OpenAI. Wenn das Unternehmen beschliesst, ChatGPT nicht mehr zur Verfügung zu stellen, dann haben wir als Nutzer:innen nichts mehr. Diese Monopolisierung zentraler Technologien ist ein antidemokratischer Vorgang par excellence. Eine dritte Form der Undurchsichtigkeit liegt in der individuellen Kompetenz der Nutzerinnen und Nutzer. Die meisten von uns sind keine Computer- und Algorithmenfachleute. Wir gebrauchen also ein Tool, das uns immer mehr gebraucht. Wir passen unsere Lebensform ihm an. Was sich bereits auf das Bildungskonzept auswirkt. So wie Sprache und Mathematik sich als kulturelle Grundkompetenzen durchsetzen, wird jetzt die Forderung nach einem gewissen «algorithmischen Alphabetismus» laut, als notwendige Anpassungsleistung an die neuen KI-Umwelten. Die schleichende Anpassung an die Geräte ist kaum aufzuhalten. Das wissen die Technounternehmen. Man halte sich nur vor Augen, wie X / Twitter den politischen Meinungsaustausch beeinflussen kann. Eine solche Technologie setzt sich in der Regel nicht demokratisch durch, sie «bürgert» sich «ein», ohne dass Bürgerinnen und Bürger ein Wort mitreden können. Wir spielen mit den Artefakten das Spiel bereits, wenn das Plebiszit ansteht, ob wir es spielen wollen.

Demokratische Risiken

Wenn ich sage, den Algorithmen sei die Demokratie scheissegal, dann gilt das erst recht für die Herren der Algorithmen. Musk sagt schon jetzt, wie der Karren läuft. Die EU hat ein Digitalgesetz ausgetüftelt, um die Plattformkonzerne in die Pflicht zu nehmen. Insbesondere ist die EU-Kommission gewillt, Musk Auflagen zu machen und ihn zu bestrafen, wenn er nicht für mehr Sicherheit bei X sorgt. Das kümmert Musk freilich wenig. Er pöbelt nur über Europa und dessen Regierungen. Er ist der Schutzgott der antidemokratischen Bewegungen. Er weiss: X steht über der Demokratie.

Die Öffentlichkeit verlagert sich zunehmend auf die Online-Plattformen. Und hier herrschen in weiten Teilen Propaganda, Desinformation, Diffamierung, Unterstellung, Lüge vor. Der Wahlkampf in den USA hat das in aller schmutzigen Klarheit gezeigt. Der ganze Unflat in der politischen Debatte ist von Algorithmen getrieben. Das Design von Plattformen wie X, Instagram und Tiktok favorisiert extreme Inhalte, die tendenziell mehr Aufmerksamkeit erregen. Viralität ist heute das ausschlaggebende Kriterium des politischen Diskurses. Wer Tabus kalkuliert bricht, kann auf die Assistenz von Algorithmen zählen. Überschreitungen der politischen Anstandsgrenzen – sprich: Pöbeleien – gehören zum neuen politischen Stil. Auch in der Schweiz. Perfid daran ist, dass gewählte Politiker auf ihre Immunität setzen können, wenn sie über die Schnur hauen. Der Anreiz, sich dadurch politisch zu profilieren, ist hoch.

Man kann über den Verfall der politischen Sitte lamentieren. Aber wie begegnet man ihm? Mit Strafnormen? Mit Sprachregelung? Wenn neue politische Sitten sich dank neuer Kommunikationstechnologien durchsetzen, soll man dann diese Technologien verbieten? Und wer definiert überhaupt «die» Sitte? Ein oberster demokratischer Sittenwächter?

Die Verwilderung des Diskurses ist offensichtlich. Und daran hat die Symbiose von Big Tech und Öffentlichkeit ihren Anteil. Gerade nach den amerikanischen Wahlen ist die Befreiung der digitalen Medien von den digitalen Monopolisten ein intensiv diskutiertes Problem. Gerade Parteien mit Interesse an Wettbewerb, Marktoffenheit und -vielfalt sollte die Loslösung geäusselter Meinungen aus den Fängen digitaler Monopolplattformen interessieren.

Gesellschaftliche Herausforderungen

Ich komme noch einmal auf die Anfangsfrage zurück: Wie können Artefakte überhaupt als «Akteure» wirken? Meine These lautet: Das «Akteurhafte» der Artefakte ist Symptom dafür, dass wir uns selbst als Akteur:innen immer mehr «herausnehmen»; kognitive Fähigkeiten nicht mehr kultivieren, die unerlässlich für das demokratische Spiel sind. Und dieses Spiel beruht nicht nur auf dem freien und vielgestaltigen Austausch der Meinungen, es braucht eine Schiedsinstanz, die für den Einhalt der Spielordnung sorgt. Und zuallererst braucht es das demokratische Grundmotiv par excellence, das sich im Grunde als erkenntnistheoretisches entpuppt: den Willen zur Wahrheit.

Der spanische Philosoph Ortega y Gasset schrieb in seinem Buch «Der Aufstand der Massen» (1929): «Wer Ideen haben will, muss zuerst die Wahrheit wollen und sich die Spielregeln aneignen, die sie auferlegt. Es geht nicht an, von Ideen oder Meinungen zu reden, wenn man keine Instanz anerkennt, welche über sie zu Gericht sitzt».

Man könnte Demokratie als Gesellschaftsform definieren, die das alltägliche Leben durch solche übergeordneten Instanzen zu regeln sucht; sei dies im Recht, sei dies in den politischen und wirtschaftlichen Beziehungen. Und eben auch im Meinungsaustausch. Man sollte sich bewusst sein, dass die digitalen Medien Meinungsverklumpung fördern, und das bedroht das fragile «Wir» einer freien, offenen, rechtsstaatlich geregelten Gesellschaft. In diesem Zusammenhang stiess ich kürzlich auf einen interessanten etymologischen Fingerzeig. Ursprünglich soll nämlich das deutsche Wort «meinen» wenig mit behaupten zu tun gehabt haben. Es ist vielmehr verwandt mit Gemeinschaft, Allgemeinheit, Gemeinsinn. Wer eine Meinung vertritt, signalisiert den Willen, einen Beitrag zu einer gemeinsamen Sache zu leisten. Kurz gesagt, bedeutet dies, Demokratie nicht ständig als selbstverständliche Errungenschaft zu zelebrieren, sondern sie als ein politisches Gut vor Augen zu halten, das es erst zu entdecken und zu erringen gilt.

Gefahr für die Wahrheit

Die sogenannten intelligenten Tools zersetzen den Willen zur Wahrheit, weil ihnen die Wahrheit scheissegal ist. Der Chatbot generiert eine Sequenz von Wörtern, und er «weiss» nicht, ob diese Sequenz wahr ist, geschweige denn, ob sie überhaupt etwas bedeutet.

Die grösste Feindin der Demokratie lauert so gesehen in der Ambivalenz der neuen Technologien: Sie erweitern unsere kognitiven Vermögen und machen zugleich träge, indem sie uns dazu verführen, unsere eigenen kognitiven Vermögen zu vernachlässigen. Wir beobachten dies ja gerade aktuell an den ChatGPTs. Sie schreiben nicht, sie simulieren Schreiben, generieren eine Sequenz von Wörtern, und sie «wissen» nicht, ob diese Sequenz wahr ist, geschweige denn, ob sie überhaupt etwas bedeutet. Wenn wir also die Frage stellen «Warum noch schreiben lernen, wenn die Maschine dies ebenso gut kann?», dann sind wir der Maschine schon auf den Leim gekrochen. Die Frage drückt eine gefährliche Resigniertheit aus, die an eine berühmte Passage aus Kants Schrift «Was ist Aufklärung?» erinnert: «Es ist so bequem, unmündig zu sein. Habe ich ein Buch, das für mich Verstand hat, einen Seelsorger, der für mich Gewissen hat, einen Arzt, der für mich die Diät beurteilt usw., so brauche ich mich ja nicht selbst zu bemühen. Ich habe nicht nötig zu denken, wenn ich nur bezahlen kann; andere werden das verdriessliche Geschäft schon für mich übernehmen.»

Demokratien geraten nicht nur «von aussen», etwa durch Machtwillen, unter Druck, sondern gerade auch «von innen», durch das Nicht-nötig-haben-zu-denken. Ich behaupte nicht, dass «intelligente» Technologie uns dumm mache. Ihr wohnt freilich die Tendenz inne, dass wir uns selbst dumm machen, indem wir meinen, den KI-Systemen das Denken zu überlassen. Sie stecken noch in den Kinderschuhen. Sie werden

«Bestimmte Fähigkeiten zu automatisieren, bedeutet nicht, auf diese Fähigkeiten zu verzichten.»

erwachsener, und sie infiltrieren künftig unseren sozialen Verkehr. Werden sie diesen Verkehr schleichend prägen – bis auch uns Menschen die Wahrheit unserer Sätze scheissegal erscheint? Technologie entlastet uns von vielem. Sobald wir uns aber vom Denken entlastet wähnen, befinden wir uns auf schiefer antidemokratischer Bahn. Die technisch avanciertesten Gesellschaftsformen sind zugleich die demokratisch prekärsten.

Die neuen «intelligenten» Geräte bestechen uns, in einem Doppelsinn des Wortes. Sie bestechen uns durch ihre teils übermenschlich anmutenden Fähigkeiten. Und sie bestechen uns mit ihrem verführerischen Versprechen, nahezu alle sozialen und politischen Probleme effizient zu lösen. Heute, im Universum der smarten Dinge, entgehen wir dieser Bestechung kaum noch. Uns aus dem Bannkreis dieser Bestechung zu lösen, wäre ein erster Schritt technologischer Aufklärung. Der scharfsichtige Philosoph und Publizist Günter Anders prägte den Begriff der Antiquiertheit des Menschen. Zunehmend übernimmt die Maschine menschliche Fähigkeiten. Und sie übertrifft ihn oft. Hüten wir uns nur vor dem Fehlschluss: Bestimmte Fähigkeiten zu automatisieren, bedeutet nicht, auf diese Fähigkeiten zu verzichten. Der Mensch ist ein äusserst erfinderisches Wesen. Er kann also in einer Welt der Technik sich neu erfinden. Ja, er kann auch Demokratie neu konzipieren. Nicht dadurch, dass er ihre Prozesse automatisiert, sondern dadurch, dass er sich zweimal überlegt, ob er sie an Geräte delegieren will. Ja, mehr noch: Wissen wir überhaupt, was wir alles können? Ist das Delegieren an Geräte ein Nullsummenspiel? Verlieren wir notwendig unsere kognitiven Vermögen?

Wir neigen dazu, alle sozialen und politischen Veränderungen dem Einfluss der Technologie zuzuschreiben, und wir vergessen dabei, dass es «die» Technologie nicht gibt. Es gibt Menschen – Ingenieure, Unternehmer, Investoren, KI-Evangelisten –, welche die Technologie zu ganz bestimmten Zwecken einsetzen – und missbrauchen. Und vielen von ihnen liegt durchaus daran, dass die Nutzer:innen ihrer Produkte in der Herdenwärme einer lammfrommen Technikgläubigkeit verharren. Das demokratische Projekt, das hier Kontur gewinnt, wäre ein anthropologisches Projekt im Sinne Kants: technikmündiger Mensch zu werden.

Dr. Eduard Kaeser ist Physiker und promovierter Philosoph.

Was ist KI und wie funktioniert sie?

Was ist Künstliche Intelligenz und denkt sie wirklich?

Die Entwicklungen in der Künstlichen Intelligenz werfen grundlegende Fragen auf: Was genau ist KI, und ist sie wirklich in der Lage zu denken? Der Philosoph **Claus Beisbart** nähert sich diesen Fragen, indem er zwischen Denken und Handeln unterscheidet. Denken bedeutet dabei die Ausführung von Algorithmen, während Handeln auf die Fähigkeiten von Robotern verweist. Ein zentrales Unterscheidungsmerkmal ist der Unterschied

zwischen schwacher und starker KI. Schwache KI simuliert menschliche Fähigkeiten ohne echtes Denken, während starke KI Maschinen anstrebt, die in der Lage sind, eigenständig zu denken und zu handeln.

Prof. Dr. Dr. Claus Beisbart
Universität Bern, Institut für Philosophie



 **Referat
auf Youtube**



 **Referat
auf Youtube**

Das Gehirn als Vorbild

Die Funktionsweise des menschlichen Gehirns dient als Grundlage für die Entwicklung von Künstlicher Intelligenz. Wie **Mihai Petrovici** vom Institut für Physiologie der Universität Bern erläutert, funktionieren Neuronen dabei wie kleine Schalter: Sie sammeln Signale und «feuern» sie weiter, wenn eine bestimmte Schwelle erreicht ist. Diese Signale fließen durch anpassbare Verbindungen (Synapsen), die durch Nutzung stärker werden und damit Lernen ermöglichen. Frühe Erwartungen an KI mit menschenähnlichen Fähigkeiten erfüllten sich oft nicht, doch mithilfe tiefer neuronaler Netzwerke und Fehlerrückpropagation

lernen künstliche Netze, indem Fehler erkannt und die Verbindungen angepasst werden. Fortschritte in Big Data und Grafikprozessoren beschleunigen diese Entwicklung, brauchen jedoch viel Energie – anders als das menschliche Gehirn, das mit etwa 20 Watt auskommt.

Dr. Mihai A. Petrovici
Universität Bern, Institut für Physiologie



Wie Maschinen lernen

Maschinelles Lernen ist die Fähigkeit, Wissen durch Erfahrungen zu erwerben und Entscheidungen auf Basis von Daten zu treffen. **Sigve Haug** vom Data Science Lab der Universität Bern beschreibt das «Agent-Umwelt»-Modell, bei dem ein Agent Informationen aus seiner Umgebung wahrnimmt, verarbeitet und zielführend handelt. Die drei Hauptarten des maschinellen Lernens umfassen überwachtes Lernen, bei dem Maschinen durch Beispiele lernen, verstärktes Lernen, das durch

Belohnung und Bestrafung die Entscheidungsfindung steuert, und unüberwachtes Lernen, bei dem Muster eigenständig erkannt werden.

PD Dr. Sigve Haug
Universität Bern, Data Science Lab

Bilder bearbeiten, Bilder manipulieren

Die Möglichkeiten der KI-gestützten Bildbearbeitung sind weitreichend. Durch Programme wie Mid-Journey lassen sich realistische Porträts generieren, analysieren und manipulieren. Neben der statischen Bildbearbeitung erlaubt KI auch Animationen: Gesichter können blinzeln oder sprechen, was in interaktiven Videos genutzt werden kann. **Petra Müller** vom Institut für Psychologie der Universität Bern betont jedoch die Risiken: KI-gestützte Manipulationen wie Deepfakes machen es schwer, zwischen realen und gefälschten Inhalten zu unterscheiden. Diese Gefahr zeigt sich besonders

in Bereichen wie Betrug und Identitätsdiebstahl, bei denen KI erstellte Bilder oder Videos für Täuschungszwecke genutzt werden könnten. Solche Entwicklungen stellen Herausforderungen für die Strafverfolgung dar, die Schwierigkeiten haben könnte, in der Masse an generierten Inhalten Relevantes zu erkennen.

Petra Müller
Universität Bern, Institut für Psychologie



 **Referat
auf Youtube**

Die dunkle Seite der Macht?

Algorithmen selektionieren politische Informationen

Künstliche Intelligenz beeinflusst die Suche nach politischen Informationen und verändert die Art, wie Menschen Wissen finden und verarbeiten. **Silke Adam** vom Institut für Kommunikations- und Medienwissenschaft der Universität Bern beleuchtet zwei zentrale Anwendungen: Algorithmen in Suchmaschinen und generative Sprachmodelle. Beide Technologien verschmelzen immer stärker, wobei Suchmaschinen wie Google nicht mehr nur Quellen auflisten, sondern direkt Antworten liefern.

Die Nutzung solcher Systeme ist weit verbreitet, das Vertrauen in sie hoch – Studien zufolge halten 75 % der Nutzerinnen und Nutzer Google für vertrauenswürdig. Dennoch mangelt es vielen am Verständnis, wie personalisierte Ergebnisse entstehen. Adam fordert mehr

Transparenz über Algorithmen und Trainingsdaten sowie eine stärkere Förderung von Medienkompetenz. Angesichts der Marktdominanz weniger Unternehmen und der ethischen Herausforderungen betont sie: KI ist nicht neutral – sie prägt unsere Wahrnehmung und beeinflusst gesellschaftliche wie politische Entscheidungen nachhaltig.

Prof. Dr. Silke Adam

Universität Bern, Institut für Kommunikations- und Medienwissenschaft



 **Referat
auf Youtube**



Eine neue Dimension der Propaganda

Deepfake-Technologien ermöglichen es, täuschend echte Video- und Audioinhalte zu erstellen, die Desinformationen verbreiten und Unsicherheit stiften können. **Jennifer Scurrrell**, Doktorandin am Center for Security Studies der ETH Zürich, spricht über die Auswirkungen Künstlicher Intelligenz auf moderne Propaganda. Sie erläutert, wie Technologien wie Deepfakes und KI-gestützte Chatbots zur Manipulation eingesetzt werden: Ein bekanntes Beispiel ist das gefälschte Video des ukrainischen Präsidenten Wolodymyr Selenskyj, der dazu aufruft, die Waffen niederzulegen. Weiter beeinflussen KI-Chatbots Meinungen, indem sie personalisierte

Botschaften erstellen und bestimmte Kommunikationsstile imitieren, was ihre Glaubwürdigkeit und Wirkung verstärkt. Dank ihrer Automatisierung und Skalierbarkeit verbreiten solche Systeme Propaganda in einer nie dagewesenen Geschwindigkeit und Reichweite.

Jennifer Victoria Scurrrell
ETH Zürich, Center for Security Studies (CSS)



 **Referat
auf Youtube**

Nutzen und Grenzen von KI

Die rasante Entwicklung von Künstlicher Intelligenz hat hohe Erwartungen geweckt, von der Automatisierung bis zur Schaffung neuer Möglichkeiten. **Ruth Fulterer**, Wirtschaftsredaktorin bei der Neuen Zürcher Zeitung, beleuchtet zwei zentrale KI-Arten: analytische KI, die Datenmuster analysiert und in Bereichen wie Diagnostik präzise arbeitet, und generative KI, die Inhalte wie Texte und Bilder erstellt. Letztere zeigt sowohl beeindruckende Ergebnisse als auch Schwächen, darunter die Erfindung falscher Informationen («Halluzinationen»). Fulterer betont, dass die Nützlichkeit von KI besonders in kreativen Aufgaben oder komplexen, leicht überprüfbaren Szenarien liegt. Gleichzeitig erfordert die Anwendung von KI Wissen und Kontrolle, da Fehler nur durch menschliches Eingreifen minimiert werden können.

Trotz beeindruckender Fortschritte bleibt KI aktuell in vielen Bereichen unzuverlässig. Grosse Technologieunternehmen wie Google, Microsoft und OpenAI dominieren den Markt, da die Entwicklung von KI immense Datenmengen, teure Hochleistungsprozessoren und hohen Energieverbrauch erfordert. Diese Ressourcen fördern die Marktkonzentration und erschweren kleinen Akteuren den Zugang. Die Balance zwischen Nutzen, Kosten und Nachhaltigkeit wird die zukünftige Entwicklung prägen.

Ruth Fulterer
Redaktorin Neue Zürcher Zeitung (NZZ)



Neuronale Abläufe mit Lu.i verstehen

Unser Gehirn besteht aus Milliarden von Neuronen, die über elektrische und chemische Signale miteinander kommunizieren und sich ständig an neue Eindrücke anpassen. Diese Fähigkeit zur **Lernfähigkeit** und die parallele Verarbeitung von Informationen machen das Gehirn aussergewöhnlich effizient. Inspiriert von diesen biologischen Prozessen sind künstliche neuronale Netze in der Künstlichen Intelligenz grob nachgebildet: Sie simulieren, wie Neuronen Signale empfangen, weiterleiten und durch veränderbare Verbindungen (Gewichtung) Muster erkennen und daraus lernen.

Trotz der grossen Bedeutung neuronaler Informationsverarbeitung gibt es nur wenige anschauliche Werkzeuge, um diese Prinzipien greifbar zu machen. Eine Ausnahme bildet **Lu.i** – ein elektronisches Modell, das von einem **Team um Mihai A. Petrovici** an der Universität Bern entwickelt wurde.

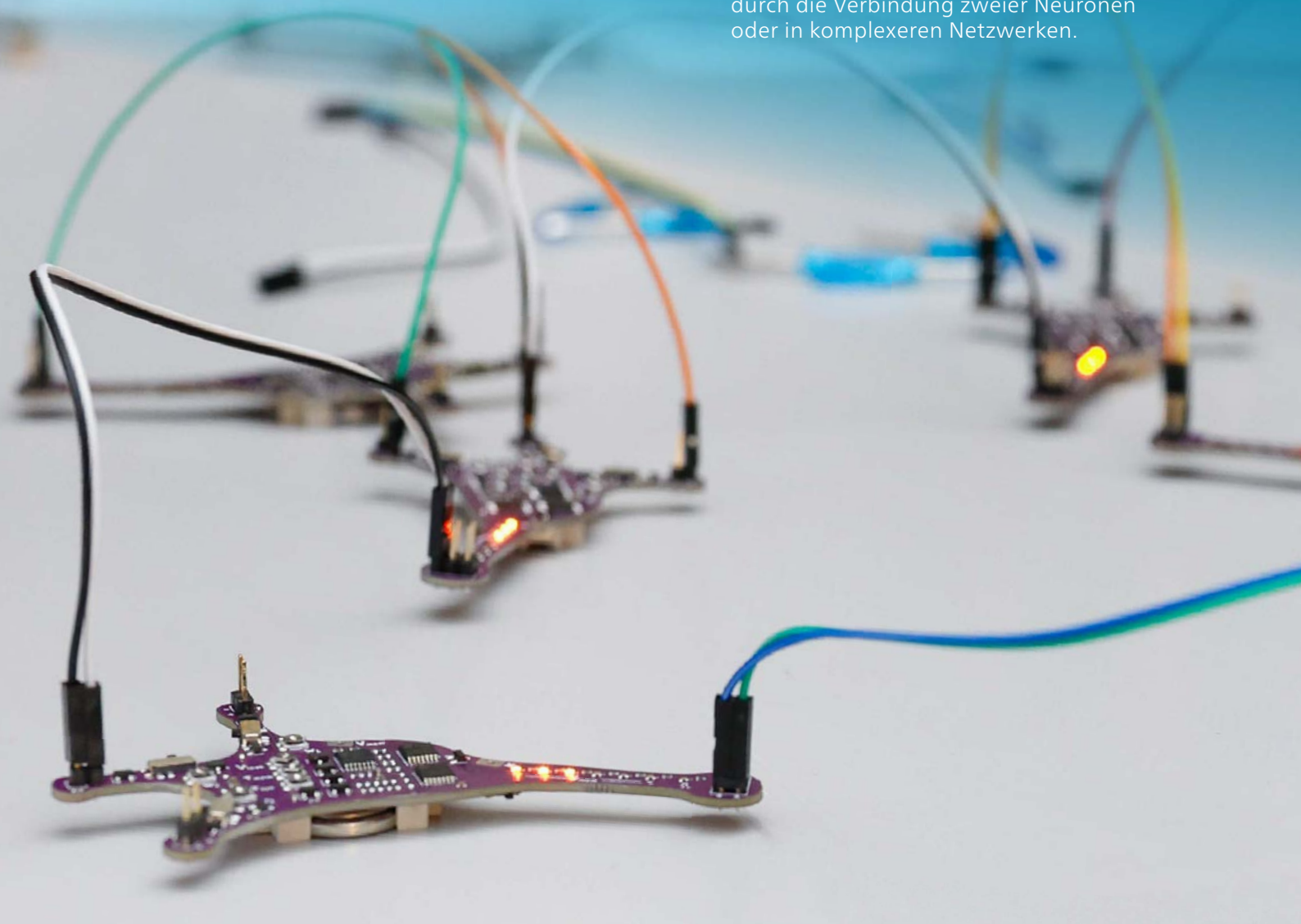
Lu.i ist ein **kleines elektronisches Neuron**, das zwei wesentliche Eigenschaften echter Nervenzellen veranschaulicht:

1. Integration von Signalen über Zeit und Raum

- Eingehende Signale werden von Lu.i aufaddiert.
- Schwache Signale «lecken» (d.h. sie bauen sich langsam ab).
- Sobald ein bestimmter Schwellenwert überschritten wird, feuert das Neuron und sendet ein Signal weiter.
- Dieser Prozess wird durch eine LED-Anzeige sichtbar gemacht, die das Feuern des Neurons signalisiert.

2. Selektive Signalweitergabe

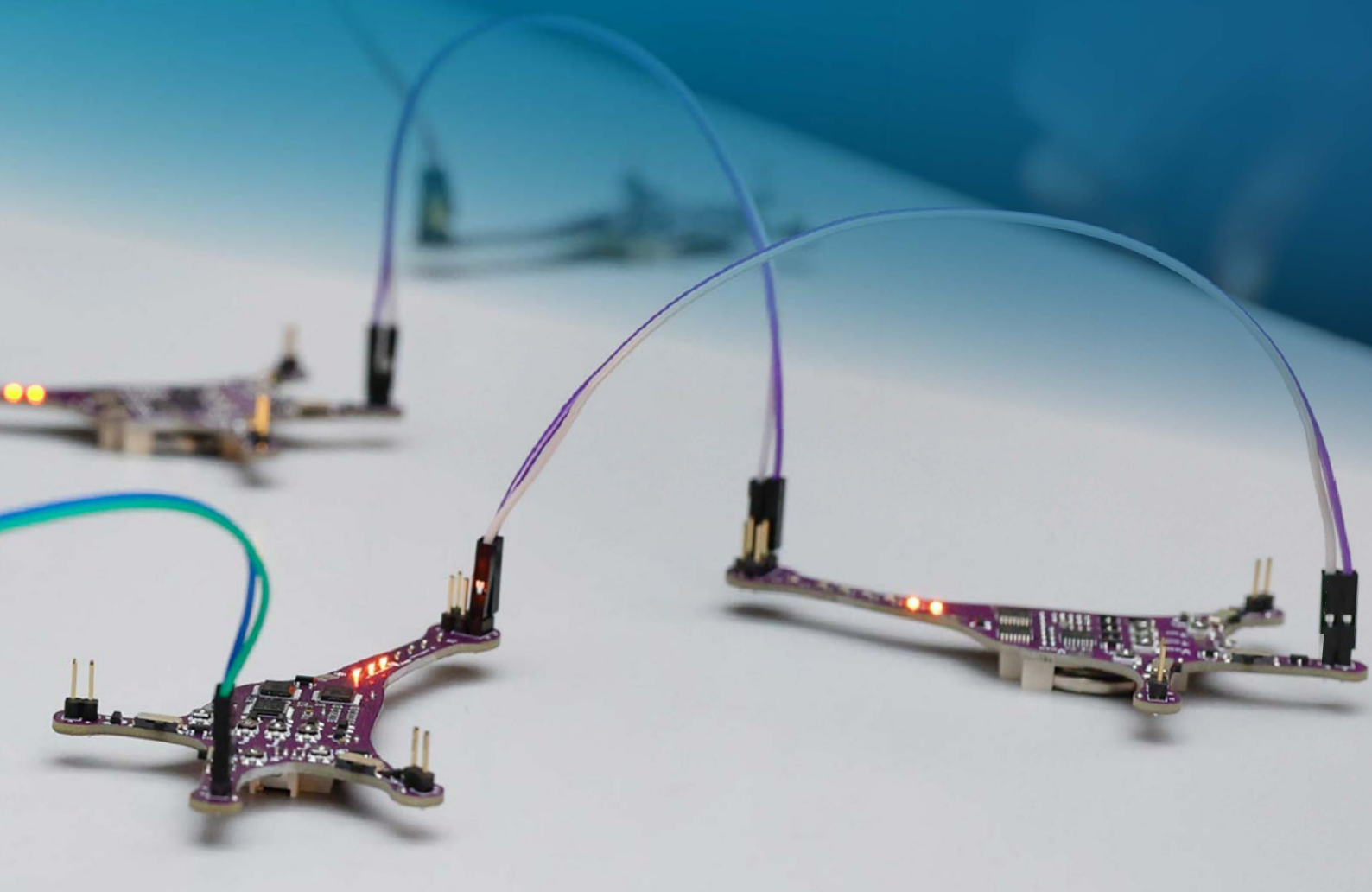
- Neuronen senden Informationen nur weiter, wenn die Eingangsreize stark genug sind.
- Dies lässt sich mit Lu.i in verschiedenen Experimenten nachvollziehen – etwa durch die Verbindung zweier Neuronen oder in komplexeren Netzwerken.



Während Lu.i die elektrische Signalverarbeitung innerhalb eines Neurons demonstriert, bildet es die chemische Signalübertragung zwischen Neuronen nicht nach. In biologischen Neuronen werden in den Synapsen **Neurotransmitter** freigesetzt, die an Rezeptoren der nächsten Nervenzelle binden und dort ein neues elektrisches Signal auslösen.

Dank Lu.i lassen sich die grundlegenden Prinzipien der neuronalen Informationsverarbeitung spielerisch und intuitiv erfassen – ein wertvolles Werkzeug für Lehre, Forschung und Wissenschaftskommunikation, wie auch die Seiten 2–3 in diesem Heft zeigen.

Marcus Moser



Künstliche Intelligenz für natürliche Gesundheit?



Krankheiten überwinden dank KI?

Proteine sind winzige biologischen Nano-maschinen und bilden die Basis aller Lebensprozesse. Sie bestehen aus 20 verschiedenen Aminosäuren, die wie Perlen an einer Kette miteinander verbunden sind. Von der Form unserer Knochen über die Kraft unserer Muskeln bis hin zur Sauerstoffübertragung durch Hämoglobin – Proteine sind an allen essenziellen Prozessen beteiligt. **Thomas Lemmin** beschreibt Proteine als Sätze in der Sprache des Lebens. Doch wie ein Tippfehler die Bedeutung eines Satzes verändern kann, kann eine Mutation in der DNA die Funktion eines Proteins stark beeinträchtigen. Dies ist besonders relevant in der medizinischen Forschung, wo solche Veränderungen oft die Ursache für Krankheiten sind. Hier kommt die Künstliche Intelligenz ins Spiel. Lemmin erklärt, wie Transformer-Modelle, wie sie

auch bei Sprach-KI wie ChatGPT verwendet werden, genutzt werden können, um die Sprache der Proteine besser zu verstehen. Diese Modelle erstellen hochdimensionale Karten, um Beziehungen und Bedeutungen zwischen Aminosäuresequenzen zu analysieren. Dadurch wird es möglich, Proteinstrukturen schneller vorherzusagen und neue Zusammenhänge zu entdecken. Die Nutzung von KI zur Modellierung von Proteinen bietet enormes Potenzial, biologische Probleme schneller und präziser zu lösen.

Ass. Prof. Dr. Thomas Lemmin
Universität Bern, Institut für Biochemie
und Molekulare Medizin

Potential und Grenzen von Chatbots in der Psychotherapie

Ein Drittel der Schweizer Bevölkerung leidet innerhalb eines Jahres an einer psychischen Erkrankung. Besonders Depressionen und Angststörungen bleiben häufig unerkannt oder unbehandelt, teils aufgrund von Scham oder fehlender Therapieplätzen. Innovative Lösungen sind dringend erforderlich, betont **Thomas Berger** vom Psychologischen Institut der Universität Bern. Digitale Technologie könnte hier Abhilfe schaffen, indem sie die Reichweite psychologischer Interventionen erhöht. Sie ist jedoch kein Ersatz für menschliche Therapie. Der persönliche Kontakt bleibt ein entscheidender Faktor für die

Wirksamkeit. Die Zukunft liegt in der Kombination beider Ansätze, um den Zugang zu psychologischer Hilfe zu verbessern und die gesellschaftlichen Auswirkungen psychischer Erkrankungen zu mindern. Gleichzeitig ist eine sorgfältige Regulierung notwendig, um die Sicherheit und Effektivität dieser Ansätze zu gewährleisten.

Prof. Dr. Thomas Berger
Universität Bern, Institut für Psychologie



Bessere Diagnostik dank KI

Künstliche Intelligenz ist aus der modernen Medizin nicht mehr wegzudenken. Insbesondere in der bildgebenden Diagnostik spielt sie bereits eine entscheidende Rolle. **Piotr Radojewski** vom Universitätsinstitut für Diagnostische und Interventionelle Neuroradiologie des Inselspitals erläutert, wie KI-basierte Systeme Radiolog:innen unterstützen und zunehmend komplexe Aufgaben übernehmen.

Ein Hauptproblem der Radiologie ist das exponentiell wachsende Bildmaterial. Hier hilft die KI, indem sie Veränderungen in MRT- und CT-Bildern automatisch erkennt und markiert. Fortschritte zeigen sich insbesondere bei Krankheiten wie Multipler Sklerose oder Demenz. Ein Beispiel verdeutlicht dies: KI erkennt subtile Veränderungen in Gehirnbildern, die für das menschliche Auge kaum wahrnehmbar sind, und liefert eine klare quantitative Analyse. Der Einsatz von KI beschränkt sich nicht nur auf die Analyse einzelner Bilder; die Systeme organisieren auch ganze Arbeitsprozesse, etwa bei der Schlaganfallbehandlung. Die KI synchronisiert die Abläufe verschiedener medizinischer Etappen und spart so wertvolle Zeit – ein entscheidender Faktor für die Patientenversorgung. Zentral bleibt jedoch die Bedeutung qualitativ hochwertiger Daten für die Wirksamkeit der Algorithmen. Fehlende oder unpassende Daten könnten die Leistungsfähigkeit der Systeme erheblich einschränken.

Dr. med. Piotr Radojewski

Inselspital Bern, Universitätsinstitut für
Diagnostische und Interventionelle Neuroradiologie



 [Referat auf Youtube](#)



Roboter zur Unterstützung älterer Menschen

Die Zahl der über 65-Jährigen wächst, und viele möchten möglichst lange in ihrer gewohnten Umgebung bleiben. Während Internet und smarte Geräte wie Staubsaugerroboter teilweise genutzt werden, gibt es Vorbehalte gegenüber Pflegerobotik und Sozialrobotern, wie **Alexander Seifert** von der Hochschule für Soziale Arbeit FHNW erläutert. Ältere Generationen haben oft Schwierigkeiten mit neuen Technologien, und viele digitale Dienste sind schwer zugänglich. Pflegeroboter wie der Roboterbär oder die emotionale Robbe «Paro» bieten Potenzial, werden aber noch wenig verwendet. Studien zeigen positive Effekte, doch ethische Fragen bleiben ungeklärt. Die Entscheidung für Technik treffen oft Angehörige,

nicht die Betroffenen selbst. Jüngere sind der Robotik gegenüber aufgeschlossener, während Ältere skeptischer sind. Digitalisierung ist unausweichlich, aber Technik sollte nur dann eingesetzt werden, wenn sie wirklich nützlich und sicher ist. Abschliessend wird von Seifert betont, dass Technik menschliche Fürsorge nicht ersetzen kann, sondern nur unterstützen soll.

Dr. Alexander Seifert

Hochschule für Soziale Arbeit FHNW,
Institut Integration und Partizipation

Ethische Implikationen beim KI-Einsatz

Stephen Milford beleuchtet in seinem Vortrag die ethischen Herausforderungen des KI-Einsatzes in der Medizin und betont die Bedeutung der Menschlichkeit. KI wie ChatGPT kann Diagnosen und Patientenkommunikation mit hoher Genauigkeit übernehmen und könnte in unterversorgten Regionen helfen. Gleichzeitig stellt sich die Frage, ob Maschinen Ärzte ersetzen sollten, da der Mensch sich durch soziale Beziehungen definiert. Besonders die Arzt-Patient-Beziehung ist essenziell für den Behandlungserfolg, da Empathie zur Heilung beiträgt. Eine ungleiche Versorgung könnte entstehen, wenn reiche Patienten mensch-

liche Betreuung erhalten, während ärmere Menschen auf KI angewiesen sind. Trotz technologischer Fortschritte müssen ethische Standards den verantwortungsvollen Einsatz gewährleisten. Medizin bedeutet mehr als nur Diagnosen – sie umfasst auch Fürsorge und Verständnis. KI sollte Ärzte unterstützen, aber die persönliche Betreuung nicht ersetzen. Letztlich bleibt die menschliche Beziehung ein zentraler Bestandteil der Heilung.

Dr. Dr. Stephen Milford
Universität Basel, Institute for Biomedical Ethics



 [Referat auf Youtube](#)



Die Zukunft?

Bildung und Regulation

KI und Kreativität – Eine Frage der Anerkennung

Kann eine KI Kunst oder Literatur erschaffen? Der Philosoph und Literaturwissenschaftler **Hannes Bajohr** argumentiert, dass die Frage nicht technisch, sondern sozial zu beantworten ist. Kunst entsteht nicht durch formale Kriterien wie Neuheit oder Originalität, sondern durch Anerkennung innerhalb eines kulturellen Rahmens.

Die Debatte über KI und Kreativität krankt laut Bajohr an einem verkürzten Verständnis des Kreativitätsbegriffs. Ursprünglich als göttliches Privileg betrachtet, wurde Kreativität in der Romantik mit dem Genie-Konzept verbunden. Heute ist sie vor allem ökonomisch aufgeladen: Die Kreativwirtschaft fordert unablässige Innovation. Die KI-Forschung greift diesen Begriff auf, reduziert ihn jedoch auf messbare Parameter und ignoriert den sozialen Kontext von Kunst.

Ein zentrales Problem bleibt: Kunst existiert nicht isoliert, sondern in einem Geflecht aus kulturellen Institutionen, Diskursen und

Anerkennungsmechanismen. Eine KI kann zwar Texte generieren, aber sie kann sich nicht als Autorin etablieren. Kunst, so Bajohr, ist nicht nur eine Frage der Produktion, sondern der sozialen Zuschreibung.

Mensch-KI-Kollaborationen zeigen, dass maschineller Output erst durch menschliche Kuratierung Bedeutung erhält. Projekte wie «Pharmaco-AI» interpretieren KI-Texte als Teil eines kreativen Unbewussten, doch auch hier bleibt der Mensch die Instanz, die entscheidet, was Kunst ist.

Die entscheidende Frage ist daher nicht, ob KI kreativ sein kann, sondern wer über den Kunststatus entscheidet. Solange dieser Anerkennungsakt menschlich bleibt, bleibt auch Kunst menschlich.

Ass. Prof. Dr. Hannes Bajohr
University of California, Berkeley



 [Referat auf Youtube](#)



 [Referat auf Youtube](#)

KI und die Universität – Vom Wissen zum Können

Künstliche Intelligenz stellt die Hochschulen vor eine grundlegende Herausforderung: Wenn Maschinen Texte produzieren, die mit denen von Studierenden konkurrieren, muss sich die universitäre Lehre neu ausrichten.

Fritz Sager, Vizerektor Lehre an der Universität Bern, plädiert für ein Umdenken – weg von der Wissensabfrage, hin zur Förderung analytischer und kreativer Fähigkeiten.

Die erste Reaktion auf ChatGPT war an der Universität oftmals Panik. Plötzlich schien die Exklusivität akademischer Fähigkeiten infrage gestellt. Doch wie einst die Schreibmaschine oder das Internet muss KI in den akademischen Alltag integriert werden. Prüfungen, die reines Faktenwissen abfragen, haben an Bedeutung verloren. Künftig geht es darum, Wissen kritisch anzuwenden.

Ein neuer Blick auf Fehler zeigt zudem den Unterschied zwischen Mensch und Maschine: Während KI Wahrscheinlichkeiten berechnet, basiert Wissenschaft auf überprüfbaren

Zusammenhängen. Maschinen liefern perfekte Texte, doch gerade Unvollkommenheit kann Ausdruck menschlicher Originalität sein.

Die Universität Bern setzt auf aktive Gestaltung dieses Wandels: Forschungsgruppen und Lehrveranstaltungen widmen sich dem sinnvollen Einsatz von KI. Entscheidend ist nicht die Kontrolle, ob ein Text von einer KI oder einem Menschen stammt, sondern die Kompetenz, KI als Werkzeug kritisch zu nutzen. Die Universität der Zukunft prüft nicht mehr nur Wissen, sondern Können.

Prof. Dr. Fritz Sager
Universität Bern, Vizerektor Lehre



 [Referat auf Youtube](#)

KI – Macht ohne Kontrolle?

Die Regulierung von Künstlicher Intelligenz bleibt eine der grössten Herausforderungen unserer Zeit. **Estelle Pannatier** von AlgorithmWatch CH warnt vor intransparenten Algorithmen, die Grundrechte unterlaufen und demokratische Strukturen gefährden können.

In der Schweiz werden KI-Systeme bereits in zahlreichen Bereichen eingesetzt – von der vorausschauenden Polizeiarbeit bis zur Stimmenauthentifizierung bei Banken. Doch oft ist unklar, wer diese Technologien entwickelt und steuert. Algorithmen beeinflussen Kreditvergaben, Jobchancen und sogar politische Meinungsbildung, ohne dass ihr Einfluss nachvollziehbar wäre.

Neben diesen gesellschaftlichen Risiken weist Pannatier auf die vernachlässigte Nachhaltigkeitsdimension hin. Die enormen Ressourcen, die KI verbraucht und die schlechten Arbeitsbedingungen in der Plattformökonomie stellen drängende Probleme dar. Die Dominanz weniger Technologieunternehmen führt zudem zu einem Machtgefälle, das demokratische Prinzipien untergräbt.

Die bestehenden Gesetze bieten bislang keine ausreichenden Antworten. Während der Bundesrat erste Schritte unternimmt, fordert Pannatier eine umfassende Strategie zur Steuerung von KI. Es gehe nicht darum, KI entweder zu verteufeln oder zu glorifizieren, sondern realistische Risiken zu adressieren. Ohne transparente Governance bleibe KI eine Blackbox – mit unkalkulierbaren Folgen für Gesellschaft und Demokratie.

Estelle Pannatier
Policy Managerin AlgorithmWatch CH

KI-Regulierung – Die Schweiz setzt auf Zurückhaltung

Während die EU mit dem AI-Act strikte Regeln schafft, wählt die Schweiz einen anderen Weg. Statt spezifischer Technologiegesetze setzt sie auf sektorale Anpassungen und internationale Kooperation, wie **Sabine Brenner** vom Bundesamt für Kommunikation (BAKOM) aufzeigt.

Das Bundesamt hat im Auftrag des Bundesrats eine Bestandsaufnahme der KI-Regulierung erstellt. Der Bericht zeigt: Bestehende Gesetze zu Datenschutz, Antidiskriminierung und Produktsicherheit erfassen bereits viele Aspekte der Technologie. Eine umfassende Regelung, die gezielt auf KI-Risiken eingeht, fehlt jedoch. Ein wichtiger Schritt ist die Unterzeichnung der KI-Konvention des Europarats, die Grundrechte im Umgang mit KI sichern soll. Parallel arbeitet das Bundesamt für Justiz an einer Vernehmlassungsvorlage für weitergehende Vorschläge bis 2026. Statt einer technologiezentrierten Regulierung – wie sie die EU verfolgt – setzt die Schweiz auf eine Wirkungsregulierung, die gesellschaftliche und rechtliche Konsequenzen in den Fokus rückt.

Branchenbezogene Massnahmen existieren bereits: Im Strassenverkehr wurden Rahmenbedingungen für autonomes Fahren geschaffen, in Zoll- und Finanzbehörden wird KI genutzt. Bis zu einer umfassenderen Regelung setzt der Bund auf freiwillige Standards und den Dialog mit Unternehmen.

In der Verwaltung steckt der KI-Einsatz noch in den Anfängen. Einzelne Anwendungen wie maschinelle Übersetzung sind etabliert, doch ein übergreifendes Konzept fehlt. Die geplante «Swiss Government Cloud» soll hier Abhilfe schaffen.

Die Schweiz verfolgt einen pragmatischen Kurs: keine überstürzte Regulierung, aber enge internationale Zusammenarbeit. Offen bleibt, ob dieser Ansatz ausreicht, um Innovation und Schutz gleichermaßen zu gewährleisten.

Sabine Brenner

Leiterin Digital Office Bundesamt für Kommunikation BAKOM



 [Referat auf Youtube](#)

KI und Wissen – Zwischen Präzision und Täuschung

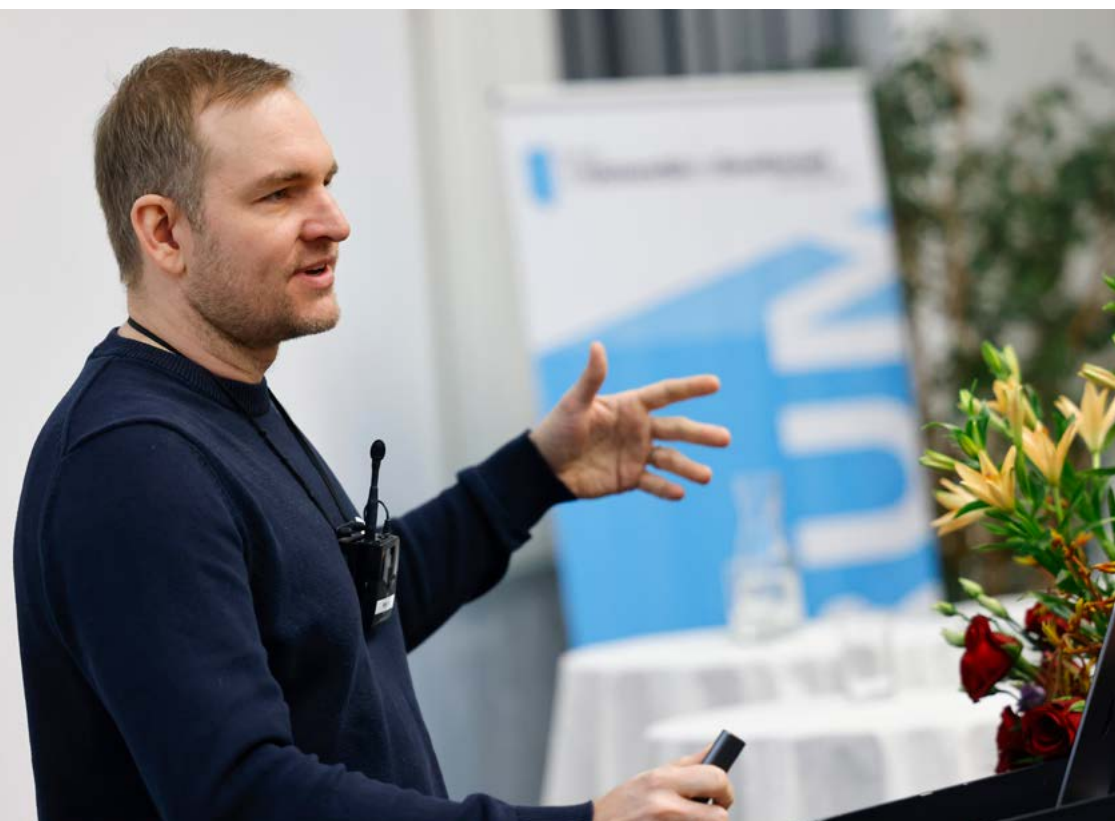
Large Language Models (LLMs) haben den Zugang zu Wissen revolutioniert. Doch sie erzeugen keine originären Erkenntnisse, sondern berechnen lediglich die wahrscheinlichste Fortsetzung eines Satzes. Das Ergebnis sind Texte, die perfekt formuliert, aber oft inhaltlich leer oder fehlerhaft sind.

Tobias Hodel von der Universität Bern warnt vor den Risiken dieser Technologie. Wer KI-generierte Inhalte unkritisch übernimmt, riskiert die Verbreitung von Pseudowissen. Besonders problematisch sind erfundene Quellenangaben und intransparente Trainingsdaten. Gleichzeitig sieht Hodel Potenzial für eine produktive Nutzung von LLMs, etwa zur Strukturierung grosser Datenmengen oder zur Generierung erster Entwürfe. Entscheidend sei jedoch, dass Nutzende die Mechanismen hinter den Modellen verstehen und deren Grenzen erkennen.

Ein weiteres Problem ist die Abhängigkeit von grossen Technologieunternehmen. Die dominierenden KI-Modelle stammen aus den USA, ihre Trainingsdaten sind nicht offengelegt. Wissenschaftliche Institutionen sollten daher eigene, spezialisierte Modelle entwickeln, um Kontrolle über die Wissensproduktion zu behalten.

Schliesslich fordert Hodel neue Prüfungsformate: Studierende müssen zeigen, dass sie argumentieren und reflektieren – nicht nur, dass sie Texte generieren lassen können. «Data Literacy», also die Fähigkeit, KI-generierte Inhalte kritisch zu analysieren, wird zur Schlüsselkompetenz der Zukunft. Denn nicht Maschinen, sondern Menschen müssen Wissen bewerten und einordnen.

Prof. Dr. Tobias Hodel
Universität Bern, Digital Humanities



 [Referat auf Youtube](#)



Künstliche Intelligenz in Anwendung

Künstliche Intelligenz findet heute Anwendung in vielen Bereichen und beeinflusst unser Leben auf vielfältige Weise. Hier sind einige markante Beispiele:

1. Sprachverarbeitung und Übersetzung

- Chatbots und virtuelle Assistenten: Systeme wie Siri, Alexa und Google Assistant verwenden KI, um Sprachbefehle zu verstehen und zu verarbeiten.
- Maschinelle Übersetzung: Plattformen wie Google Translate und DeepL nutzen KI, um Texte und gesprochene Sprache in Echtzeit in andere Sprachen zu übersetzen.
- Texterstellung und -zusammenfassung: KI kann Artikel schreiben, Texte zusammenfassen und sogar kreative Geschichten oder Gedichte verfassen.

2. Bild- und Videoverarbeitung

- Gesichtserkennung: Wird in Sicherheitssystemen, von Social-Media-Plattformen und auch zur Authentifizierung auf Smartphones verwendet.
- Bildklassifizierung und -analyse: KI kann Objekte, Szenen oder Gesichter in Bildern erkennen und klassifizieren (z.B. in der Medizin zur Erkennung von Tumoren auf Röntgenbildern).
- Bildverbesserung und -restauration: Algorithmen können Bildqualität verbessern, Details rekonstruieren und sogar alte Fotos oder Videos farblich aufbereiten.

3. Empfehlungssysteme

- Personalisierte Inhalte: Plattformen wie Netflix, YouTube und Spotify verwenden KI, um Empfehlungen basierend auf bisherigen Aktivitäten zu geben.
- E-Commerce: Amazon und andere Online-Shops nutzen KI, um Kunden Produkte zu empfehlen, die sie interessieren könnten.
- Soziale Medien: Facebook, Instagram und Twitter setzen Algorithmen ein, die den Feed personalisieren, indem sie aufgrund des Nutzungsverhaltens die am meisten relevanten Inhalte anzeigen.

4. Gesundheitswesen

- Diagnose und Bildgebung: KI kann Anomalien in medizinischen Bildern wie MRTs, CT-Scans oder Röntgenbildern erkennen.
- Medikamentenentwicklung: KI hilft bei der Analyse und Identifikation neuer potenzieller Wirkstoffe, was die Entwicklung von Medikamenten beschleunigt.
- Überwachung und Vorhersage: Wearables nutzen KI, um Gesundheitsdaten in Echtzeit zu überwachen und mögliche Gesundheitsrisiken vorherzusagen.

5. Autonomes Fahren

- Autonome Fahrzeuge: Unternehmen wie Tesla, Waymo und Uber arbeiten an selbstfahrenden Autos, die mithilfe von KI Umgebungen erkennen und Navigationsentscheidungen treffen können.
- Assistenzsysteme: Viele moderne Autos nutzen bereits KI für Fahrerassistenzsysteme, die Spur halten, bremsen oder beschleunigen können.

6. Finanzwesen und Wirtschaft

- Algorithmischer Handel: Banken und Hedgefonds verwenden KI, um den Handel automatisch durchzuführen und Marktmuster zu erkennen.
- Betrugserkennung: Kreditkartenunternehmen und Banken nutzen KI, um verdächtige Transaktionen zu identifizieren und Betrug zu verhindern.
- Kundendienst: Chatbots und automatische Kundensupportsysteme nutzen KI, um häufig gestellte Fragen zu beantworten und Kunden zu unterstützen.

7. Industrie und Produktion

- Qualitätskontrolle: KI-gestützte Systeme analysieren Produktionsprozesse und erkennen Fehler schneller und genauer als Menschen.
- Robotik: Roboter in Fabriken sind oft mit KI ausgestattet, um komplexe Aufgaben autonom durchzuführen, wie etwa das Sortieren oder Verpacken.
- Wartung und Überwachung: KI-gestützte Systeme prognostizieren den Wartungsbedarf,

um Maschinenausfälle zu minimieren und die Produktion zu optimieren.

8. Bildung

- Personalisierte Lernplattformen: Systeme wie Duolingo und andere E-Learning-Plattformen passen den Lerninhalt an das Niveau und das Tempo des Benutzers an.
- Automatische Bewertung: KI kann schriftliche Arbeiten und Tests bewerten und sofortiges Feedback geben.
- Virtuelle Tutoren: Einige KI-gestützte Systeme bieten individuelle Nachhilfe und Unterstützung für Schüler.

9. Landwirtschaft

- Erntevorhersagen und Bodenanalyse: KI analysiert Wetterdaten, Bodenbeschaffenheit und andere Faktoren, um die besten Anbaumethoden und Erntezeiten zu bestimmen.
- Robotik und Automatisierung: Autonome Maschinen pflanzen, düngen und ernten Nutzpflanzen.
- Schädlingsüberwachung: KI erkennt Muster von Schädlingsbefall und unterstützt bei der Wahl effektiver Schädlingsbekämpfungsstrategien.

10. Kreative Bereiche

- Kunst und Musik: KI-generierte Kunstwerke und Musikstücke werden immer beliebter und auch als Inspirationsquelle verwendet.
- Design und Architektur: KI hilft beim Entwerfen von Gebäuden oder Produkten und optimiert Designentscheidungen basierend auf funktionalen und ästhetischen Faktoren.
- Schreiben und Content Creation: Einige Algorithmen können eigenständig Artikel schreiben, visuelle Effekte in Videos erstellen und sogar Texte für Werbung verfassen.

Häufig verwendete Begriffe kurz erklärt

Algorithmus

Eine Reihe von Regeln oder Schritten, die ein Computer befolgt, um ein bestimmtes Problem zu lösen oder eine Aufgabe durchzuführen.

Artificial General Intelligence (AGI)

Die Idee einer KI, die über umfassende, menschenähnliche Intelligenz verfügt und flexibel verschiedene Aufgaben lösen kann.

Automation

Der Einsatz von KI, um menschliche Aufgaben zu übernehmen, oft zur Effizienzsteigerung, z.B. in der Industrie oder im Kundenservice.

Bias (Verzerrung)

Eine systematische Voreingenommenheit in KI-Modellen, die auf unausgewogenen oder fehlerhaften Trainingsdaten basiert und zu unfairen oder diskriminierenden Ergebnissen führen kann.

Big Data

Grosse Datenmengen, die zur Analyse und Entscheidungsfindung verwendet werden und das maschinelle Lernen ermöglichen.

Bot

Ein Bot ist ein automatisiertes Programm, das Aufgaben eigenständig ausführt. Bots können einfache Prozesse wie Datenabrufe oder komplexere Interaktionen, wie bei Chatbots, übernehmen. Sie werden oft für Kundenservice, Datensammlung oder Automatisierung eingesetzt.

Chatbot

Ein KI-gestütztes Programm, das mit Menschen interagiert und Fragen beantworten oder Aufgaben erledigen kann.

Computer Vision

Ein Bereich der KI, der es Computern ermöglicht, Bilder und Videos zu analysieren und zu interpretieren, wie z.B. bei Gesichtserkennung.

Deep Learning

Eine spezielle Form des maschinellen Lernens, bei der tiefe neuronale Netzwerke verwendet werden, die viele Schichten haben, um sehr komplexe Datenverhältnisse zu verarbeiten.

Deepfake

Ein mediales KI-Produkt, das täuschend echte Bilder oder Videos erstellt und manipuliert, oft zur Unterhaltung oder zu schädlichen Zwecken.

Explainable AI (XAI)

Ansätze, die darauf abzielen, KI-Entscheidungen transparent und für den Menschen nachvollziehbar zu machen.

KI-Ethik

Der Bereich, der sich mit den moralischen und gesellschaftlichen Auswirkungen von KI befasst, z.B. Datenschutz, Verantwortung und Fairness.

Künstliche Intelligenz (KI)

Die Fähigkeit von Maschinen, Aufgaben zu übernehmen, die normalerweise menschliche Intelligenz erfordern, wie Lernen, Problemlösen und Entscheidungsfindung.

Lernen, überwacht

Ein Lernansatz, bei dem das Modell mit gekennzeichneten Daten trainiert wird, bei denen die Eingaben und die gewünschten Ausgaben bekannt sind.

Lernen, unüberwacht

Ein Lernansatz, bei dem das Modell mit unbeschrifteten Daten trainiert wird und Muster oder Strukturen selbstständig erkennen muss.

Lernen, bestärkend (Reinforcement Learning)

Ein Ansatz, bei dem ein Modell durch Belohnungen und Strafen lernt, die besten Entscheidungen in einer bestimmten Umgebung zu treffen.

Maschinelles Lernen (ML)

Ein Teilbereich der KI, in dem Maschinen durch Daten trainiert werden, Muster zu erkennen und Vorhersagen zu treffen, ohne explizit programmiert zu sein.

Neuronales Netzwerk

Ein Modell, das dem menschlichen Gehirn nachempfunden ist und aus Schichten von «Neuronen» besteht. Es wird häufig im maschinellen Lernen eingesetzt, um komplexe Muster zu erkennen.

Natural Language Processing (NLP)

Ein Bereich der KI, der sich auf die Interaktion zwischen Computern und menschlicher Sprache konzentriert, z.B. bei Übersetzungen und Sprachassistenten.

Robotik

Der Bereich, in dem KI in physische Maschinen, sogenannte Roboter, integriert wird, die Aufgaben autonom ausführen.

Turing-Test

Ein Test, der misst, ob eine Maschine menschliches Verhalten so gut nachahmen kann, dass ein Mensch sie nicht mehr von einem Menschen unterscheiden kann.

Referent:innen

Silke Adam, Prof. Dr., ist Direktorin des Instituts für Kommunikations- und Medienwissenschaft der Universität Bern. Sie studierte Kommunikationswissenschaft an der Universität Hohenheim (D) und erwarb zudem einen Master of Science in Mass Communication an der Boston University, MA (USA). Forschungsaufenthalte führten sie an die Freie Universität Berlin (D), ans Netherland Institute of Advanced Studies (NL), an die University of Washington (USA) und an die Bournemouth University (UK). In ihrer Forschung beschäftigt sie sich mit politischer Kommunikation und deren Wandel durch Digitalisierung und Transnationalisierung. Die Verbreitung von Verschwörungsideen, die Online-Kommunikation von Klimawandelskeptikern und die selektive Medienzuwendung durch Rechtspopulisten stehen im Mittelpunkt ihrer derzeitigen Forschung.

Hannes Bajohr, Ass. Prof. Dr., ist Philosoph, Schriftsteller und Literaturwissenschaftler. Er studierte Philosophie, deutsche Literatur und neuere und neueste Geschichte an der Humboldt-Universität und wurde an der Columbia University, New York, mit einer Dissertation über Hans Blumenbergs Sprachtheorie promoviert. In seinen Werken experimentiert er häufig mit Technologie, unter anderem mit Künstlicher Intelligenz. Er lebte und arbeitete in New York, Berlin und Basel und ist inzwischen Assistenzprofessor für Germanistik an der University of California in Berkeley, USA.

Claus Beisbart, Prof. Dr. Dr., ist seit 2012 ausserordentlicher Professor für Wissenschaftsphilosophie an der Universität Bern. Er arbeitet am Institut für Philosophie, aber auch an mehreren Zentren, u.a. dem Center for Artificial Intelligence in Medicine, und er ist im Direktorium des Multidisciplinary Center for Infectious Diseases. Seit 2022 ist Claus Beisbart Präsident der Schweizerischen Philosophischen Gesellschaft. Seine jüngste Buchpublikation hat den Titel «Was heisst hier noch real? Wirklichkeiten in Zeiten von Computersimulation und virtueller Realität» (Stuttgart 2024).

Thomas Berger, Prof. Dr., leitet die Abteilung Klinische Psychologie und Psychotherapie an der Universität

Bern. Nach seinem Studium der Psychologie an der Universität Bern hat er an der Universität Freiburg i.Br. promoviert, eine psychotherapeutische Weiterbildung absolviert und in verschiedenen Institutionen klinisch gearbeitet. Er hat an der Universität Bern habilitiert. Für seine Arbeit im Bereich digitaler Interventionen in der Psychotherapie hat er verschiedene Preise erhalten, darunter 2021 den Schweizer Wissenschaftspreis Marcel Benoist.

Sabine Brenner beschäftigt sich seit 25 Jahren in verschiedenen Funktionen mit Fragen der Digitalisierung in der Bundesverwaltung. Sie leitet heute das Digital Office im Bundesamt für Kommunikation und ist Stellvertretende Stabschefin dieses Amtes. Sie ist Ko-Projektleiterin für die Erstellung eines Grundlagenberichts zuhanden des Bundesrates zum Thema Regulierung von Künstlicher Intelligenz in der Schweiz und befasst sich auch mit Anwendungsfragen von KI im Bund.

Ruth Fulterer schreibt seit 2021 als Redaktorin für Technologie für die Neue Zürcher Zeitung über Themen an der Schnittstelle zwischen Technologie und Gesellschaft. Ihr Schwerpunkt ist Künstliche Intelligenz. Zuvor war sie als freie Journalistin tätig, unter anderem für die ZEIT, das Wirtschaftsmagazin brand eins und die Süddeutsche Zeitung. Sie stammt ursprünglich aus Südtirol, Italien, und hat Volkswirtschaft und Philosophie in Wien und New Orleans studiert.

Sigve Haug, PD Dr., studierte Physik in Deutschland, Spanien und Norwegen. Er war an unterschiedlichen Experimenten der Neutrinophysik und der Hochenergiephysik beteiligt, die sich oft durch sehr grosse Datenmengen auszeichnen. Ein Beispiel ist der Large Hadron Collider am CERN in Genf. Heute leitet er das Data Science Lab an der Universität Bern, das universitätsweite Unterstützung und Schulungen im Bereich Künstlicher Intelligenz und Data Science für die datenintensive Forschung anbietet.

Tobias Hodel, Prof. Dr., ist seit 2019 Assistenzprofessor für digitale Geisteswissenschaften an der Universität Bern. Er forscht und lehrt zu maschinellen Lernverfahren in und für die

Geisteswissenschaften. Darunter fällt die automatisierte Erkennung historischer Handschriften, die Extraktion von Informationen oder die Entwicklung von spezifischen und grossen Sprachmodellen. Hodel ist promovierter Historiker und leitet an der Universität Bern unter anderem Forschungsprojekte zu den Turmbüchern der Stadt Bern aus der Frühen Neuzeit und Chatsystemen für die Hochschuldidaktik im 21. Jahrhundert.

Eduard Käser, Dr., studierte an der Universität Bern erst theoretischen Physik, anschliessend Wissenschaftsgeschichte und Philosophie. Er promovierte an der naturwissenschaftlichen Fakultät in Philosophie. Bis 2012 war er Gymnasiallehrer für Physik und Mathematik. Eduard Käser war immer wieder publizistisch tätig über Themen zwischen Wissenschaft und Philosophie, in Büchern, Magazinen und Zeitungen.

In neuerer Zeit konzentriert sich sein Interesse auf das Thema der Anthropologie im Zeitalter des Künstlichen. Zuletzt erschien «Auf schiefer Bahn» Politische Essays zur Zukunft, Basel 2023.

Thomas Lemmin, Ass. Prof. Dr., ist Assistenzprofessor am Institut für Biochemie und Molekulare Medizin (IBMM) der Universität Bern. Seine Forschung überbrückt die Lücke zwischen computergestützten und experimentellen Methoden, um komplexe biologische Systeme auf molekularer Ebene zu verstehen. Er setzt modernste Deep-Learning-Techniken ein, um die «Sprache» der Proteine zu entschlüsseln, und kombiniert diese mit traditionellen Modellierungsmethoden, um neuartige Protein-Design-Strategien zu entwickeln, die zu neuen therapeutischen Strategien führen sollen.

Stephen Milford, Dr. Dr., promovierte 2017 an der Protestantischen Theologischen Universität in den Niederlanden in den Bereichen Anthropologie, Menschenwürde, Wert, relationale Ontologie und postliberale Studien. 2024 schloss er eine zweite Promotion in biomedizinischer Ethik mit Schwerpunkt auf KI im Gesundheitswesen ab. Seine Forschungsinteressen liegen in Fragen der Identität, der Menschenrechte und der Ethik, wobei der Schwerpunkt auf relationaler

Ontologie, der Einzigartigkeit des Menschen und seiner Unersetzbarkeit liegt.

Petra Müller ist seit 2021 Doktorandin der Psychologie (UniDistance Suisse/Universität Bern) und erforscht, wie Künstliche Intelligenz menschliches Verhalten und Entscheidungsprozesse beeinflusst. Als Mitglied von Chaostreff Bern engagiert sie sich seit 2018 für digitale Rechte und ethische Technologieentwicklung. Ihr Hauptinteresse gilt der Rolle von Maschinellen Lernen in gesellschaftlichen und demokratischen Prozessen, insbesondere im Zusammenhang mit Fake News, irrationalen Überzeugungen und Verschwörungstheorien.

Estelle Pannatier ist Policy & Advocacy Managerin bei Algorithm-Watch CH. Sie hat einen Master in politischer Anthropologie und in Kommunikations- und Medienwissenschaften. Vor ihrer Tätigkeit bei AlgorithmWatch CH hat sie in der Politikentwicklung im Kontext der Digitalisierung des Bildungswesens in der Schweiz mitgewirkt. Davor hatte sie für die Online-Wahlhilfeplattform smartvote, für das Eidgenössische Departement für auswärtige Angelegenheiten EDA sowie für das Schweizer Radio und Fernsehen gearbeitet.

Mihai A. Petrovici, Dr., hat an der Universität Heidelberg in Physik promoviert und leitet seit 2020 die Forschungsgruppe «Neuro-Inspired Theory, Modeling and Application» an der Universität Bern. Diese Forschung setzt sich zwei eng miteinander verwandte Ziele: Einerseits ein besseres Verständnis des menschlichen Gehirns und seiner Fähigkeit, komplexe Aufgaben schnell und effizient zu lösen. Andererseits ermöglichen diese Erkenntnisse unmittelbare Anwendungen im Bereich der Künstlichen Intelligenz, etwa für besonders energieeffiziente Algorithmen und Mikrochips.

Piotr Radojewski, Dr. med., ist Geschäftsführer des Translational Imaging Centers am SITEM-Insel. Er ist am Institut für Diagnostische und Interventionelle Neuroradiologie des Inselspitals wissenschaftlich und klinisch tätig. Seine Verantwortungsbereiche umfassen die Koordination von Forschungsprojekten sowie die

Zusammenarbeit mit der Industrie. Er bringt umfassende Expertise an der Schnittstelle zwischen klinischer Medizin, Forschung und Industrie mit. In seiner täglichen Arbeit setzt er KI-basierte Systeme in der neuro-radiologischen Diagnostik sowie in Forschungsprojekten ein.

Fritz Sager, Prof. Dr., ist Professor für Politikwissenschaft am Kompetenzzentrum für Public Management der Universität Bern. Seine Forschungsgebiete sind die öffentliche Politik, ihre Umsetzung und ihre Wirkung, die politische Rolle der öffentlichen Verwaltung sowie das Verhältnis von Wissenschaft und Politik. Als Vizerektor Lehre der Universität befasst er sich mit der Innovation des Präsenzunterrichts mit neuen Lehr- und Lernformen unter gezieltem Einsatz digitaler Lösungen und dem Einbezug neuer Erkenntnisse aus der Lehr- und Lernforschung.

Jennifer Victoria Scurrall ist Doktorandin am Center for Security Studies (CSS) an der ETH Zürich und Red Teamerin im Bereich generative Künstliche Intelligenz (KI) für ein KI-Technologieunternehmen. Sie forscht zum Thema Human-AI Interaction und analysiert in ihrer Doktorarbeit den Einsatz von KI in Informationsoperationen mit einem Fokus auf die Beeinflussungsversuche politischer Meinungsbildungsprozesse in sozialen Netzwerken. 2023 absolvierte sie ein Forschungspraktikum bei Microsoft Research in Redmond, USA. Dort beschäftigte sie sich u.a. mit dem verantwortungsvollen Einsatz von KI.

Alexander Seifert, Dr., ist Sozialarbeiter und promovierte Soziologe. Er war 14 Jahre lang als wissenschaftlicher Mitarbeiter und Bereichsleiter Forschung am Zentrum für Gerontologie (ZfG) der Universität Zürich tätig. Seit 2020 ist er an der Hochschule für Soziale Arbeit der Fachhochschule Nordwestschweiz FHNW angestellt. Seifert hat sich auf die Alterssoziologie spezialisiert und befasst sich mit Themen wie der Techniknutzung im Alter, dem Wohnen und der Nachbarschaft älterer Menschen, der Bildung im Alter sowie den Sinnesbeeinträchtigungen und sozialen Ungleichheiten im Alter.

Moderation

Lisa Stalder ist freischaffende Journalistin, Texterin und Moderatorin und hat an der Universität Bern studiert. Während 15 Jahren war sie bei der Berner Tageszeitung «Der Bund» in verschiedenen Funktionen tätig. Am häufigsten schreibt sie über Umweltthemen, Verkehrs- und Raumplanung sowie über alles, was mit Energie zu tun hat. Mit Technik befasst sie sich auch als Mitglied des Beirats Entsorgung, der das Eidgenössische Departement für Umwelt, Verkehr, Energie und Kommunikation BAKOM im Prozess rund um Suche und Bau eines geeigneten Tiefenendlagers für radioaktive Abfälle unterstützt.

Projektgruppe

Silke Adam, Prof. Dr.
Universität Bern, Institut für Kommunikations- und Medienwissenschaft

Claus Beisbart, Prof. Dr. Dr.
Universität Bern, Institut für Philosophie

Sarah Beyeler, Dr.
Forum für Universität und Gesellschaft

Tobias Hodel, Prof. Dr.
Universität Bern, Digital Humanities

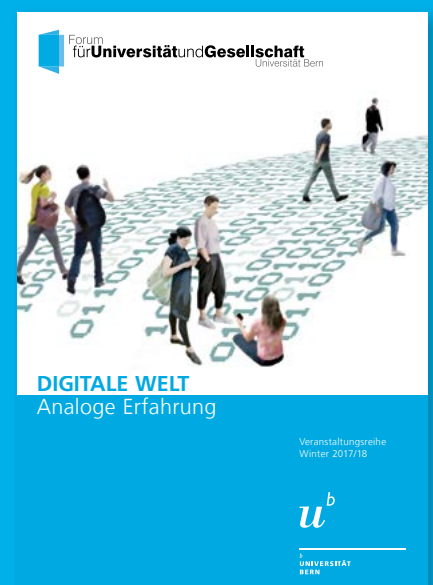
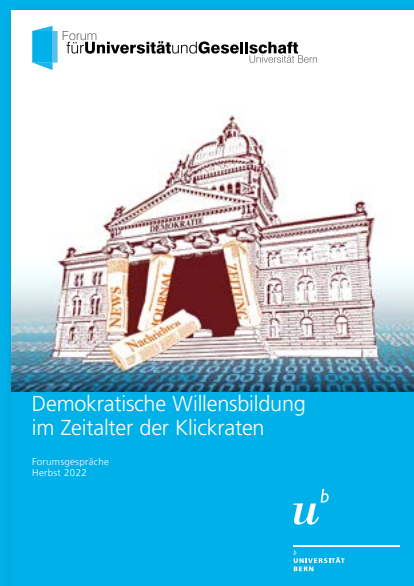
André Holenstein, Prof. em. Dr.
Universität Bern, Historisches Institut

Marcus Moser, Lic. phil.
Geschäftsleiter, Forum für Universität und Gesellschaft

Eliane Ruef
Universität Bern, Institut für Systematische Theologie

Tobias Rohrbach, Dr.
Universität Bern, Institut für Kommunikations- und Medienwissenschaft

Interessiert an den
Aktivitäten des
Forums für Universität
und Gesellschaft?
Besuchen Sie unsere Webseite:
www.forum.unibe.ch



Themenhefte vermitteln einen Überblick über die jeweiligen Forumsveranstaltungen.
Sie sind online auf unserer Webseite abrufbar: www.forum.unibe.ch/publikationen



Impressum

Themenheft 2025 / 12. Ausgabe

Herausgeber

Forum für Universität und Gesellschaft
Präsident: Prof. em. Dr. Heinzpeter Znoj
Geschäftsleiter: Lic. phil. Marcus Moser

Redaktion

Dr. Sarah Beyeler (sarah.beyeler@unibe.ch)
Lic. phil. Marcus Moser (marcus.moser@unibe.ch)

Bildnachweise

Titelbild: © 2024 Marcus Moser. Dieses Bild wurde mit Hilfe von DALL-E erstellt. / Seiten 2–4, 10–19, 27–31, 37: © Stefan Wermuth / Seiten 20–25: © Adrian Moser / Seite 7: © Eduard Kaeser, zvg / Seite 26: © Hannes Bajohr, zvg / Seite 38 oben: © iStock / Seite 38 unten: © 2025 Marcus Moser. Dieses Bild wurde mit Hilfe von DALL-E erstellt.

Layout

Christa Heinzer

Geschäftsstelle

Forum für Universität und Gesellschaft
Hochschulstrasse 6
3012 Bern
Telefon +41 31 684 45 66
info.fug@unibe.ch
www.forum.unibe.ch

Elektronische Ausgabe

© alle Rechte vorbehalten
ISSN 2624-568X

Mit freundlicher Unterstützung der
Stiftung Universität und Gesellschaft

Spendenkonto der Stiftung Universität und Gesellschaft:
CH39 0079 0042 9374 8157 5

Forumsgespräche 2025

Mentale Gesundheit: Herausforderungen, Ressourcen und Perspektiven

20. August, 3. und 16. September 2025

Eintritt frei, Anmeldung und Details unter www.forum.unibe.ch/psyche



Forum

für **Universität und Gesellschaft**

Universität Bern

Veranstaltungsreihe 2025/26

Krisen, Konflikte, Kriege – Die Weltordnung im Umbruch

1. und 22. November 2025,
17. Januar 2026

Eintritt frei, Anmeldung und Details unter www.forum.unibe.ch/weltordnung