



Hochschulforum
Digitalisierung

ARBEITSPAPIER NR. 86 / MÄRZ 2025

Künstliche Intelligenz: Grundlagen für das Handeln in der Hochschullehre

Ergebnisse der Arbeitsgruppe „Künstliche Intelligenz: Essenzielle Kompetenzen an Hochschulen“

Alexander Filipović, Aljoscha Burchardt, Simon David Hirsbrunner,
Antje Michel, Anna Puzio, Gabi Reinmann, Philipp Schaumann, Anja-Lisa
Schroll, Ulrike Tippe, Martin Wan, Nicolaus Wilder

Arbeitspapier Nr. 86 / März 2025

Künstliche Intelligenz: Grundlagen für das Handeln in der Hochschullehre

Ergebnisse der Arbeitsgruppe „Künstliche Intelligenz: Essenzielle
Kompetenzen an Hochschulen“

Autor:innen:

Alexander Filipović
Aljoscha Burchardt
Simon David Hirsbrunner
Antje Michel
Anna Puzio
Gabi Reinmann
Philipp Schaumann
Anja-Lisa Schroll
Ulrike Tippe
Martin Wan
Nicolaus Wilder

Das Hochschulforum Digitalisierung

Als bundesweiter Think and Do Tank führt das Hochschulforum Digitalisierung (HFD) eine breite Community rund um die digitale Transformation an Hochschulen zusammen, macht Entwicklungen sichtbar und erprobt innovative Lösungsansätze. Dazu werden Akteur:innen aus den Feldern Hochschulen, Politik, Wirtschaft und Gesellschaft vernetzt.

Das 2014 gegründete Hochschulforum Digitalisierung ist eine gemeinsame Initiative des Stifterverbandes, des CHE Centrum für Hochschulentwicklung und der Hochschulrektorenkonferenz (HRK). Gefördert wird es vom Bundesministerium für Bildung und Forschung (BMBF)

Inhaltsverzeichnis

1	Einleitung.....	5
2	Die Werte und Ziele von Hochschulbildung	8
2.1	Was sind die Ziele von Hochschulbildung?	8
2.2	Welche Werte liegen den Zielen von Hochschulbildung zugrunde?	8
3	Durch KI veränderte Wertvorstellungen und Zielsetzungen	11
3.1	Szenario I: (Akademische) Integrität	12
3.2	Szenario II: Autonomie	22
3.3	Szenario III: Forschung	25
4	Handlungsperspektiven für die Hochschulen im Kontext von KI: Möglichkeiten der Kompetenzentwicklung.....	30
4.1	Übergeordnete Handlungsempfehlungen	30
5	Literatur	39

1 Einleitung

Seit Tools aus dem Bereich der **generativen Künstlichen Intelligenz (KI)** mit der Veröffentlichung von ChatGPT im November 2022 als ernstzunehmende Entwicklungen auch einer breiteren Öffentlichkeit bekannt geworden sind und sich deren Entwicklung in den folgenden Monaten noch einmal deutlich beschleunigt hat, zeichnen sich auch für den Bereich der Hochschulen immense Auswirkungen ab. Dabei sind politische, ökonomische, soziale und bildungstheoretische Dimensionen erkennbar.

Bereits vielfach diskutiert werden Fragen nach der Herkunft von Informationen, ihrer Verlässlichkeit, ihrer Tendenz und Parteilichkeit sowie ihrer Zitationsfähigkeit. Auch mit KI verbundene Probleme des Datenschutzes, der Datensicherheit und des Urheberrechts drängen mit großer Vehemenz auf die Agenda. Neben den **Auswirkungen auf Studium, Lehre und Prüfungen** muss zudem auch die teils unbeabsichtigte, teils gezielte Ausbeutung und damit Entwertung wissenschaftlicher Arbeit berücksichtigt werden.

Bislang noch unzureichend beleuchtet wurde demgegenüber die Frage, wie mit beabsichtigten und unbeabsichtigten Manipulationen von Daten und in der Folge von Wissen umgegangen werden kann. Problematisch sind sowohl die **impliziten, systemimmanenten Biases** abhängig von konkreten Sprachmodellen bzw. der ihnen zugrundeliegenden Datenbasis als auch bewusste Verfälschungen wie gezielte Propaganda oder die KI-Chatbots immanente Illusion eines sozialen/menschlichen Gegenübers: Da mittels natürlicher Sprache mit Chatbots interagiert wird, sind wir mit dem Phänomen konfrontiert, dass der Mensch-Maschine-Interaktion tatsächlich leicht Eigenschaften der Mensch-Mensch-Interaktion zugeschrieben werden können. Das wirft einerseits **Fragen des sozialen Miteinanders** auf und macht andererseits eine Debatte zu genuin menschlichen Fähigkeiten und technikhärenten Grenzen Künstlicher Intelligenz notwendig.

Diesen Fragestellungen kommt über den hochschulischen Rahmen hinaus eine hohe und weiter wachsende **Bedeutung für eine Vielzahl gesellschaftlicher Kontexte** zu. Dabei hängt die Entwicklung längst nicht mehr an einzelnen Tools oder Applikationen, sondern hat mittlerweile ein raumgreifendes digitales Ökosystem entstehen lassen. In dem Maße, in dem (generative) Künstliche Intelligenz sämtliche Lebensbereiche durchdringt, gewinnt die Kenntnis ihrer grundsätzlichen Funktionsweise, insbesondere über damit verbundene Wertschöpfungsketten, an Bedeutung. So wird ein umfassendes **Verständnis von Gestalt, Wirkmechanismen und wesentlichen Treibern** dieser digitalen Ökosysteme zur Voraussetzung ihrer adäquaten Nutzung, aktiven Gestaltung und Weiterentwicklung sowie einer wissenschaftlich kritischen Auseinandersetzung mit ihnen.

Aufgabe von Hochschulen ist es, ihre Absolventen zu befähigen, als mündige Bürger in dieser Welt zu agieren. Sie stehen damit insbesondere vor der Herausforderung, ihren Angehörigen Fähigkeiten zu vermitteln, die es ihnen ermöglichen, in einer zunehmend von KI-Systemen geprägten Umwelt informierte Entscheidungen zu treffen, um so **gesellschaftliche Teilhabe auch in Zukunft** zu ermöglichen. Dazu gehört auch das Urteilsvermögen, wo KI sinnvoll eingesetzt werden kann und für welche Anwendungen und Aufgaben sie nicht geeignet ist. Denn nur wenn alle Hochschulangehörigen unabhängig von ihrem Fachgebiet die Grundprinzipien und Grenzen von KI verstehen, können Fehleinschätzungen verhindert und ihre Potenziale angemessen genutzt werden, was zu einer konstruktiven gesellschaft-

lichen Auseinandersetzung mit neuen Technologien beiträgt.

Von Dezember 2023 bis Dezember 2024 hat sich die **AG „Künstliche Intelligenz: Essenzielle Kompetenzen an Hochschulen“** mit diesen Fragestellungen beschäftigt. Dabei ging die Zielsetzung der Arbeitsgruppe über praktische Anwendungshinweise, Handlungsempfehlungen oder Praxisbeispiele zu einzelnen, aktuellen Applikationen wie etwa ChatGPT hinaus, da sich diese angesichts der Dynamik und der zu erwartenden Vielfalt der Entwicklungen in diesem Bearbeitungsformat schnell zu überleben drohen.

Vielmehr ging es im Rahmen klassischer Think-Tank Arbeit um eine grundsätzliche Auseinandersetzung mit der Frage, wie an den Hochschulen ein **kompetenter, reflektierter und wissenschaftlich fundierter Umgang mit KI** institutionell verankert und konkret implementiert werden kann. Ziel dieser Betrachtungen sollte die Ermittlung essentieller Kompetenzen im Umgang mit KI für Hochschulangehörige auf allen Ebenen sein.

In einem ersten Schritt wurde die Frage erörtert, welche Ziele Hochschulbildung im Blick hat. Daraus abgeleitet wurden in einem zweiten Schritt die Werte, die diesen Zielen von Hochschulbildung zuzuordnen sind. Mit Blick auf Künstliche Intelligenz wurde dann drittens untersucht, inwiefern sich diese **Werte und Ziele durch KI** verändern. Aus diesen Vorüberlegungen wurden schließlich im letzten Schritt Handlungsperspektiven für die Hochschulen im Kontext KI sowie Möglichkeiten der Kompetenzentwicklung erarbeitet.

Die AG hat in **Form von Szenarien** gearbeitet, um diese Überlegungen auf eine praktisch-anschauliche Ebene zu heben. Die Szenarien beschäftigen sich insbesondere mit den Themen **Integrität, Autonomie und Forschung**. Gleichzeitig erheben diese Szenarien in einem sich dynamisch entwickelnden Feld wie der Künstlichen Intelligenz keinerlei Anspruch auf eine schwerlich zu treffende abschließende Beobachtung, sondern laden zur Reflexion ein.

Den Vorsitz der AG hatte Prof. Dr. Alexander Filipović (Lehrstuhlinhaber Christliche Sozialethik, Katholisch-Theologische Fakultät der Universität Wien).

Mitglieder der AG waren:

- Dr. Aljoscha Burchardt: Stv. Standortsprecher Deutsches Forschungszentrum für Künstliche Intelligenz / KI-Campus
- Dr. Simon Hirsbrunner: Teamleiter, Internationales Zentrum für Ethik in den Wissenschaften, Universität Tübingen
- Prof. Dr. Antje Michel: Professorin für Mediendidaktik und Wissenstransfer, FH Potsdam
- Dr. Anna Puzio: Ethics of Socially Disruptive Technologies (ESDiT), Universität Twente
- Prof. Dr. Gabi Reinmann: Hochschuldidaktikerin, Universität Hamburg
- Dr. Philipp Schaumann: Vorsitzender KMK-AG Künstliche Intelligenz / Ministerium für Wissenschaft und Kultur Niedersachsen

- Prof. Dr. Ulrike Tippe: Präsidentin TH Wildau / HRK-Vizepräsidentin für Digitalisierung und wissenschaftliche Weiterbildung
- Dr. Nicolaus Wilder: Wissenschaftlicher Mitarbeiter, Institut für Pädagogik, CAU Kiel

Betreut wurde die AG durch Martin Wan, Anja-Lisa Schroll und Stella Berendes von der Geschäftsstelle der Hochschulrektorenkonferenz in Bonn im Rahmen des Hochschulforums Digitalisierung.

2 Die Werte und Ziele von Hochschulbildung

2.1 Was sind die Ziele von Hochschulbildung?

Hochschulen haben nach Huber (1983)¹ drei Kernfunktionen². Sie sind:

- Bildungseinrichtung: Sie bieten umfassende Qualifikation und dienen somit der **Praxis**.
- Wissenschaftssystem: Sie generieren neues **Wissen** und fördern den wissenschaftlichen Nachwuchs.
- Lebens- und Arbeitswelt: Sie ermöglichen **Personen** Bildung und Forschungsfreiraum.

Diese drei Funktionen – bezogen auf Praxis, Wissenschaft und Person – bilden den Bezugspunkt der Universität als Institution und ihrer Lehre. 2015 hat der Wissenschaftsrat diese Trias als Ziele der Hochschulbildung formuliert:

*„Drei zentrale Dimensionen spannen den Raum hochschulischer Bildungsziele auf: **(Fach-)Wissenschaft, Persönlichkeitsbildung und Arbeitsmarktvorbereitung**. Die Dimension (Fach-)Wissenschaft wird insbesondere von Qualifizierungszielen bestimmt, die darauf ausgerichtet sind, die Studierenden zur situationsgerechten Auswahl, Anwendung und Anpassung wissenschaftlicher Methoden sowie zum selbständigen und kritischen Umgang mit wissenschaftlichen Erkenntnissen zu befähigen. Ziele, die die Entwicklung einer fachlichen Identität sowie eines wissenschaftlichen und beruflichen Ethos befördern oder auf eine Verantwortungsübernahme im Beruf und im gesellschaftlichen Leben vorbereiten sollen, können vor allem der Dimension Persönlichkeitsbildung zugeordnet werden. Die Dimension Arbeitsmarktvorbereitung betrifft schließlich die Qualifizierung der Studierenden, die unmittelbar und gezielt auf das Erwerbsleben nach dem Studienabschluss – innerhalb oder außerhalb der Wissenschaft – ausgerichtet ist.“³*

2.2 Welche Werte liegen den Zielen von Hochschulbildung zugrunde?

Hochschulen verfolgen nicht nur hochschulische Bildungsziele im engeren Sinne, sondern

1 Huber, Ludwig, Hochschuldidaktik als Theorie der Bildung und Ausbildung, in: Ders. (Hg.), Ausbildung und Sozialisation in der Hochschule (= Enzyklopädie Erziehungswissenschaften 10), Stuttgart 1983, 127-129.

2 Vgl. Reinmann, Gabi; Watanabe, Alice, KI in der universitären Lehre, in: Schreiber, Gerhard; Ohly, Lukas (Hg.), KI: Text. Diskurse über KI-Textgeneratoren, Berlin 2024, 29-46, 31 f.

3 Wissenschaftsrat, Empfehlungen zum Verhältnis von Hochschulbildung und Arbeitsmarkt. Zweiter Teil der Empfehlung zur Qualifizierung von Fachkräften (Drs. 4925-15), Bielefeld 2015, abgerufen am 21. Jan. 2025 unter: <https://www.wissenschaftsrat.de/download/archiv/4925-15.pdf>, 9.

erfüllen auch eine Vielzahl gesellschaftlicher Funktionen. So können Hochschulen beispielsweise dazu beitragen, den sozialen Zusammenhalt zu fördern und die freiheitlich-demokratische Grundordnung zu sichern. Ein Beispiel dafür ist die „Einheit von Forschung und Lehre“, wie sie in der Humboldtschen Idee europäischer Universitäten konstitutiv ist. Den übergeordneten Zielen von Hochschulen und Hochschulbildung können **Werte** zugeordnet werden, die mit diesen Zielperspektiven eng verbunden sind.

- **Wissenschaft:** Als zentraler Leitwert steht die „wissenschaftliche Integrität“ als Ausdruck für das wissenschaftliche Ethos im Zentrum. Mit diesem Oberbegriff ist gemeint: „respektvolle[r] Umgang miteinander, mit Studienteilnehmerinnen und -teilnehmern, Tieren, Kulturgütern und Umwelt“ sowie Verantwortung angesichts der Freiheit der Wissenschaft,⁴ was auch soziale Verantwortung mit einschließt. Dies wird wiederum erreicht durch Werte, die den Forschungsprozess prägen. Zu nennen sind Qualität oder Solidität wissenschaftlicher Prozesse, etwa die Einhaltung fachspezifischer Standards und etablierter Methoden, klare Verantwortlichkeiten für Prozesse des wissenschaftlichen Arbeitens, Einbeziehung vorliegender Forschungsleistungen (Forschungsstand), Einhaltung von rechtlichen und ethischen Vorgaben (Compliance), umfassende Publikation aller Forschungsergebnisse und relevanter Information (Transparenz, Öffentlichkeitsprinzip), Anerkennung von Autor:innenschaften und Urheber:innen, Vertraulichkeit und Neutralität bei Begutachtungen.⁵ Fokussiert man die spezifischen Bildungsaufgaben von Hochschulen im Kontext der Wissenschaft, so sind Mündigkeit und Kritikfähigkeit im Umgang mit wissenschaftlichen Erkenntnissen zentrale Werte.
- **Persönlichkeitsbildung:** Die Bildung von Personen bezieht sich heute vor allem auf Autonomie im Sinne individueller Selbstbestimmung. Deci und Ryan haben in ihrer Selbstbestimmungstheorie drei psychologische Grundbedürfnisse herausgearbeitet und über mehrere Jahrzehnte in empirischen Studien belegt, dass diese sehr stabil zu finden sind: Kompetenzerleben, Autonomieerleben und soziale Eingebundenheit. Ein Studium sollte in diesem Sinne Studierenden Kompetenz- und Autonomieerleben ermöglichen und Potenziale und Erfahrungen sozialer Eingebundenheit bieten. Weiterhin ist davon auszugehen, dass reale Möglichkeiten von Partizipation in der Institution, Transparenz und Fairness Voraussetzung für eine Bindungserfahrung sind, die am Ende verantwortungsfähige Subjekte hervorbringt, die in der modernen Demokratie und in globalen Zusammenhängen gestaltungsfähig und gestaltungswillig sind. Angesichts einer Vielfalt von Kulturen, Lebensweisen und Vorstellungen des guten Lebens sind für den gesellschaftlichen Zusammenhalt Werte wie Respekt und Toleranz mit dem Ziel der Persönlichkeitsbildung verbunden.⁶

4 Deutsche Forschungsgemeinschaft, Leitlinien zur Sicherung guter wissenschaftlicher Praxis. Kodex, Bonn 2022 (korrigierte Version 1.1), abgerufen am 21. Jan. 2025 unter: https://zenodo.org/records/6472827/files/kodex_leitlinien_gwp_dfg.1.1.pdf, 7.

5 Vgl. ebd.

6 Vgl. Deci, Edward L.; Ryan, Richard M., Self-determination theory: A macrotheory of human motivation, development, and health, in: Canadian Psychology 49 (2008) Nr. 3, 182-185.

- **Arbeitsmarktvorbereitung:** Werte, die mit der Arbeitsmarktausrichtung von Hochschulen verbunden sind, sind Beteiligungsmöglichkeiten an gesellschaftlichen Prozessen durch Arbeit für Studierende sowie als gesellschaftliche Institution auch Fairness im Sinne von gerechten Chancen auf dem Arbeitsmarkt für Absolvent:innen. Diese Beteiligungsmöglichkeiten bedeuten auch, dass Studierende durch eine gute Arbeitsmarktausrichtung ihrer Ausbildung in der Lage sind, durch ihre Arbeit Gesellschaft auch zu gestalten (Mitgestaltung, Mündigkeit) und Verantwortung zu übernehmen.

3 Durch KI veränderte Wertvorstellungen und Zielsetzungen

Nach der Identifikation dieser Wert- und Zielsetzungen stellt sich die Frage, inwieweit diese durch Künstliche Intelligenz herausgefordert bzw. verändert werden:

- **Werden einzelne Werte oder Ziele von Hochschulbildung vor dem Hintergrund technischer Neuerungen umgedeutet?**
- **Verändern sich im Kontext von KI-Technologien die Realisierungsbedingungen von Zielen und Werten der Hochschulbildung?**
- **Wie können Ziele und Werte von Hochschulbildung durch verantwortungsvollen KI-Einsatz gestärkt werden?**

Diesen Fragen hat sich die AG über Szenarien angenähert, die beispielhaft für den Hochschulalltag stehen. Die Szenarien haben wir in Abwandlung bestehender Techniken der **Zukunfts- und Designforschung**⁷ in vier Schritten entwickelt: Zunächst haben wir mit allen Angehörigen der AG die oben dargestellten Werte der Hochschulbildung unterschiedlichen Situationen von Lehre und Forschung zugeordnet, Treiber und Hemmnisse identifiziert und jeweils ein grundsätzliches Spannungsfeld anhand der Expertise der Think-Tank-Mitglieder herausgearbeitet. Diese wurden von einem Zweiterteam zu Problemstellungen zugespitzt, die als Ausgangspunkt der **Szenarienentwicklung** dienten. In fachlich einschlägigen Kleingruppen wurden daraus Szenarien ausgearbeitet, deren Zielstellung im Möglichkeitsraum eines konstruktiv gefärbten wahrscheinlichen Entwicklungsszenario angelegt wurden. Die ausgearbeiteten Szenarien wurden mithilfe eines zweistufigen Peer-Review Prozesses durch die übrigen Think-Tank-Mitglieder und anschließend durch Teilnehmende eines Workshops auf dem University:Future Festival überarbeitet und validiert. Dabei wurden die Szenarien bewusst als faktische und eher positive Möglichkeitsräume entwickelt; es wurde also das Ziel verfolgt, unseren **Zugang zu Zielen und Werten in Zukünften** mit einer produktiven Nutzung von KI-Tools nachvollziehbar zu machen. Bei der Entwicklung der Szenarien orientierten wir uns an einer sehr nahen Zukunft. Teilweise realisierten sich die Rahmenbedingungen unserer Szenarien im Verlauf der gut einjährigen AG-Arbeit bereits in gegenwärtigen Entwicklungen.

Die Szenarien erfüllten in unserer Think-Tank-Arbeit und auch in diesem Dokument einen zweifachen Zweck: Zum einen spielten wir im Rahmen der Szenarien Möglichkeiten der Auseinandersetzung mit den oben beschriebenen Werten und Zielen durch, um ihre Relevanz

7 Vgl. exemplarisch: Rosson, Mary Beth; Carrol John M., Scenario-based design, in: Sears, Andrew; Jacko, Julie A. (Hg.), The Human-Computer Interaction Handbook, New York 22008, 1041-1060; De Jouvenel, Hugues, A brief methodological guide to scenario building, in: Technical Forecasting and Social Change 65 (2000) Nr. 1, 37-48; Nathan, Lisa P. u.a., Value scenarios. A technique for envisioning systemic effects of new technologies, in: CHI '07 Extended Abstracts on Human Factors in Computing Systems, San Jose 2007, 2585-2590, zu finden unter: <https://dl.acm.org/doi/10.1145/1240866.1241046> (abgerufen am 17. Feb. 2025).

für die KI-integrierende Hochschule zu durchdenken und daraus Handlungsempfehlungen für den Umgang mit KI zu entwickeln. Die Szenarien bildeten die **kondensierte Diskursgrundlage** für den Austausch mit Expert:innen, mit denen wir Methode und Ergebnisse unseres Think Tanks validierten und anreicherten. Zum anderen haben wir im Think Tank mit der Szenariotechnik einen methodischen Zugang zur Entwicklung von möglichen Zukunftsperspektiven aus Empirie und Reflexion der Gegenwart erprobt. Somit haben wir die Fähigkeit eingeübt, die mit dem Begriff der **Futures Literacy** beschrieben wird⁸ – nicht zu verwechseln mit dem im Kontext von KI gegenwärtig viel diskutierten Kompetenz-Set vermeintlicher Future Skills, das als notwendig erachtet wird, um in einer zukünftigen Berufs- und Lebenswelt erfolgreich zu bestehen.⁹ Dem methodischen Vorgehen zugrunde liegt die Annahme, dass jede Bewertung von Entwicklungen stets eine implizite Vorstellung möglicher Folgen enthält. Diese, in Bezug auf technische Innovation häufig entweder überzogen euphorisch oder dystopisch gefärbten, unausgesprochenen Vorstellungen haben wir somit expliziert und für die gemeinsame Auseinandersetzung als Orientierungsrahmen nutzbar gemacht.¹⁰ Eine begründete Vision der Zukunft, denn nichts anderes ist dieser **kollaborativ entwickelte Orientierungsrahmen**, ist für Entscheidungen in der Gegenwart essentiell, da sie Handlungsspielräume und -grenzen bewusst macht. Wir haben uns entschieden, diesen doppelten Zweck auch in diesem Bericht anzulegen und laden dazu ein, die folgenden Szenarien zum einen als Illustration der abstrakten Ziele und Werte in jeweils möglichen Zukünften wahrzunehmen. Zum anderen möchten wir mit unseren Beispielen dazu anregen, kollaborativ entwickelte Zukunftsvorstellungen als Ausgangspunkt zu nutzen, um sich zu gegenwärtigen Entwicklungen zu positionieren. Dazu lassen sich unsere Szenarien nachnutzen oder eigene Szenarien entwerfen.

3.1 Szenario I: (Akademische) Integrität

Das folgende Szenario schildert eine typische Situation, die sich so oder so ähnlich an vielen Hochschulen derzeit abspielen kann. Im Kern geht es um verschiedene Aspekte akademischer Integrität und individueller wie auch **kollektive Verantwortung in der Hochschulbildung**, die sich beim Thema Prüfungen besonders deutlich zeigen. Das Szenario liefert keine fertigen „Lösungen“ zum Umgang mit KI im Kontext Hochschulprüfungen; vielmehr lädt es zur Diskussion ein. Entsprechend steuern wir nach der Geschichte unseren eigenen Dialog bei, der sich aus dem Nachdenken über drei Leitfragen entwickelte, die sich die Expertengruppe (AG) selbst gegeben hat – in der Annahme, dass diese Leitfragen den aktuellen Diskussionsbedarf an Hochschulen widerspiegeln.

8 Vgl. Miller, Riel, Futures Literacy. A hybrid strategic scenario method, in: Futures 39 (2007) Nr. 4, 341-662.

9 Vgl. exemplarisch Ehlers, Ulf-Daniel, Future Skills. Lernen der Zukunft – Hochschule der Zukunft, Wiesbaden 2020.

10 Vgl. Miller, Riel, Transforming the future. Anticipation in the 21st century, Paris 2018.

Ausgangslage

Sarah Kaya ist an ihrer Universität Professorin für Wirtschaftsinformatik und organisiert jedes Semester ein Seminar zu Robotik. Im vergangenen Semester hat sie die Erfahrung gemacht, dass mehrere Studierende größere **Teile ihrer Hausarbeiten mit ChatGPT** geschrieben haben. Sie möchte den Einsatz generativer KI nicht einfach verbieten, sondern einen produktiven und kritischen Umgang damit fördern und einen Weg finden, die Veranstaltung so zu gestalten, dass die Studierenden Kompetenzen für einen konstruktiven Umgang mit KI im Studium entwickeln.

Konstruktive Vision

Beim Nachdenken über den Umgang mit den Risiken von ChatGPT in Prüfungssituationen wird Sarah klar, dass sie zunächst ihre bisherigen Routinen im Seminaralltag auf den Prüfstand stellen muss: **Was will ich genau mit der Hausarbeit?** Ist die Hausarbeit geeignet, die Kompetenzen aufzubauen und nachzuweisen, die in der Veranstaltung zu Robotik zentral sind? Wäre eine alternative Form des Leistungsnachweises nicht grundsätzlich besser? Die Studierenden könnten in der Vorbereitung generative KI nutzen, dann aber in Präsenz und ohne technische Hilfe ihre Arbeitsergebnisse vorstellen. Sie selbst könnte dann gleich direktes Feedback geben.

In der Cafeteria trifft Sarah ihren Kollegen Yago Wagner, wissenschaftlicher Mitarbeiter aus der Politikwissenschaft. Sara nutzt die zufällige Begegnung, um von ihrer Idee mit der neuen Prüfungsform zu berichten. Schnell wird die Debatte hitzig: Yago hält Sarahs Lösung auf die Hausarbeit zu verzichten für kurzschlüssig: Ist das nicht eine bloße Reaktion auf externe technische Entwicklungen? Sarah verteidigt ihre Entscheidung, ist aber am Ende froh über den kontroversen Austausch, denn: Das hilft ihr, sich noch klarer zu werden, warum sich in ihrem Fall eine mündliche Präsentation der studentischen Arbeiten als Prüfungsleistung fachlich sogar besser begründen lässt als eine Hausarbeit. Gleichzeitig wird ihr im Gespräch mit Yago aber auch bewusst, dass ihr Umgang mit dem GPT-Problem in Hausarbeiten nicht für alle Fächer ein Modell sein kann.

Yago hat nämlich auch etwas zum Thema zu berichten, das im Moment viele umtreibt. Er plant ein Lehrexperiment: Die Studierenden wählen selbst, ob sie die Hausarbeit ohne generative KI schreiben oder generative KI nutzen und den Einsatz dann dokumentieren und kommentieren. Die Fakultät für Sozialwissenschaften hat zu Letzterem bereits Leitlinien zur Verfügung gestellt. Vorangehen soll in der Veranstaltung eine Diskussion mit den Studierenden darüber, welche Bedeutung welche Form von Schreibkompetenz haben kann: für den späteren Beruf, für die Entwicklung der eigenen Persönlichkeit, für das Fach. Sarah ist einerseits skeptisch: Kommt das Problem bei der ersten Variante nicht durch die Hintertür wieder in den Seminarraum? Yago kontert: Geht bei Sarah an der Stelle das Misstrauen gegenüber den Studierenden durch? Andererseits ist Sarah beeindruckt: Der Umgang mit ChatGPT wird mit Yagos Vorgehen reflexiver. Zu schnell ist die Kaffeepause zu Ende; das Besprochene aber klingt bei beiden nach.

Sarah setzt ihre Idee im aktuellen Semester um. Sie erreicht in ihrer Fakultät eine **Anpassung der Prüfungsordnung**, die nun um die Option der mündlichen Präsentation erweitert ist. Das

erhöht auch künftig den Spielraum für prüfungsdidaktische Entscheidungen. Rechtzeitig zum Semester hat die Universität außerdem eine datensichere Infrastruktur für den Einsatz generativer KI zur Verfügung gestellt. Damit ist sichergestellt, dass keine Person im Seminar benachteiligt ist, wenn sie die Nutzung geeigneter Systeme im Semesterprozess erlaubt und nun auch begleitet.

Yago realisiert sein Lehrexperiment ebenfalls, hat aber zuvor mit einem großen Problem zu kämpfen: Aus dem Kreis der Studierenden wie auch der Lehrenden seiner Fakultät wird der Einwand geltend gemacht, dass die alternativ wählbaren Prüfungsformen nicht vergleichbar und daher prüfungsrechtlich bedenklich seien. Yago gelingt es, seine Frustration angesichts dieser konfligierenden Werte nicht Oberhand gewinnen zu lassen. Zusammen mit zwei Juristinnen aus der Rechtsabteilung wird eine Lösung speziell für Lehrexperimente gefunden. Diese Lösung begeistert sogar die Universitätsleitung: Sie spricht sich hochschulöffentlich dafür aus, den Herausforderungen angesichts generativer KI künftig mutiger und mit experimenteller Haltung entgegenzutreten – auch in der Lehre.

Erwägung des Szenarios anhand von Leitfragen

Werden einzelne Werte oder Ziele von Hochschulbildung vor dem Hintergrund technischer Neuerungen umgedeutet?

Was sagt unser Szenario dazu? In unserem Szenario bleiben die Ziele von Hochschulbildung, wie sie z. B. der Wissenschaftsrat [2015] als (Fach-)Wissenschaftlichkeit, Arbeitsmarktvorbereitung und Persönlichkeitsbildung formuliert hat,¹¹ auch vor dem Hintergrund von KI (implizit) bestehen: So hat die Wirtschaftsinformatikerin Sarah primär die Anforderungen im Blick, die man an Absolvent:innen mit Blick auf künftige Berufstätigkeiten stellt, wozu heute in jedem Fall Wissen und Können im Umgang mit KI gehören. Die Argumentation des Politikwissenschaftlers Yago zielt darauf ab, die Studierenden zum Nachdenken zu bringen und Verantwortung für eigene Entscheidungen zu übernehmen, womit er vor allem auf die Persönlichkeitsbildung abstellt. Beide berücksichtigen, dass KI aus der Wissenschaft nicht mehr verschwinden und künftig Teil der Fachwissenschaftlichkeit sein wird, und schließen in der Folge aus, KI für ihre Prüfungen zu verbieten. Stattdessen streben sie Prüfungsformen an, die auf unterschiedliche Weise sicherstellen, dass **KI nicht missbräuchlich eingesetzt wird**. Man findet entsprechend auch eine ganze Reihe von Äußerungen oder Andeutungen im Szenario, die darauf hinweisen, dass bisherige Werte beibehalten werden: neben der titelgebenden Integrität vor allem Verantwortung, Transparenz, Autor:innenschaft, Selbstbestimmung, Partizipation, Chancengerechtigkeit und Offenheit bzw. Experimentierfreudigkeit.

Wie ist die Leitfrage zu verstehen? Leitfrage 1 ist genau genommen in zweifacher Weise lesbar: als empirische Frage in dem Sinne, dass man **via Befragungen oder Inhaltsanalysen** prüfen könnte, ob und wie sich an Hochschulen bisherige Ziele und Werte verschieben oder ändern, also z. B. unwichtiger werden, Priorität erhalten, wegfallen, neu hinzukommen etc., oder als philosophische Frage, ob und aus welchen Gründen sich an Hochschulen bisherige

¹¹ Vgl. Wissenschaftsrat, Empfehlungen, 9.

Ziele und Werte verschieben oder ändern sollten. Grundsätzlich erscheinen uns beide Lesarten und die damit zusammenhängenden Fragen wichtig: Werden derzeit einzelne Ziele und Werte bereits umgedeutet und wird das in Zukunft wahrscheinlich weiter oder nicht mehr so sein? Wäre es wünschenswert, einzelne Ziele und Werte aktuell und in Zukunft umzudeuten oder nicht und welche Gründe lassen sich dafür anführen?

Eine empirische Beschreibung des Ist-Zustands – als Antwort auf die erste Frage – gibt eine Orientierung, wie Menschen an Hochschulen derzeit mit KI umgehen oder umzugehen denken. Das zu wissen, ist unter anderem wichtig, wenn Initiativen oder Maßnahmen zum KI-Einsatz gestaltet werden sollen, was Akzeptanz und eine gewisse Anschlussfähigkeit an das Bestehende (in der Regel) voraussetzt. Aus dem Ist-Zustand lässt sich für die Zukunft gegebenenfalls folgern, was wahrscheinlich ist. Im Rahmen unserer Auseinandersetzung mit KI in der AG stellt sich die empirische Frage aus zwei Gründen jedoch nicht: zum einen, weil dazu nicht die (zeitlichen und finanziellen) Ressourcen vorliegen, zum anderen, weil wir uns dazu entschieden haben, uns nicht (weiter) von der technologischen normativen Faktizität treiben zu lassen, sondern die **aktive Rolle zurückzugewinnen** und selbst Ziele und Werte auszuhandeln, um dann darüber nachzudenken, wie KI als Mittel eingesetzt werden kann.

Eine philosophische **Bestimmung des Soll-Zustands** eröffnet den Raum, die Zukunft als eine mögliche und gestaltbare wahrzunehmen, und läuft darauf hinaus, auszuleuchten, was für den Menschen wünschenswert ist. Es gilt, zu bestimmen, ob wir in Bezug auf Ziele und Werte Prämissen haben (und warum), die durch technische Neuerungen nicht tangiert werden. Im Zuge unserer Diskussion zu KI in der AG inklusive der Entwicklung des vorliegenden Szenarios haben wir bereits Werte gesetzt und in die Geschichte eingebettet. Innerhalb des Szenarios stellt sich also zunächst einmal gar nicht die Frage, ob Ziele und Werte umgedeutet werden (sollen), sondern es werden allenfalls Thesen sichtbar zur Frage ihrer **Relevanz und Umsetzbarkeit unter KI-Bedingungen** (was uns zur zweiten Leitfrage führt). Mit anderen Worten: Mit dem Szenario haben wir die philosophische Frage aus unserer Perspektive vorläufig beantwortet, verstehen diese Antwort aber nicht als eine absolute, sondern als eine mögliche; diese soll zugleich ein Impuls sein, sie in der Hochschulöffentlichkeit zu diskutieren und dabei verschiedene Perspektiven abzuwägen, unter denen man auch zu anderen Antworten kommen könnte.

Was folgern wir daraus? Letztlich plädieren wir nach Abwägung unseres Szenarios mit Blick auf die erste Leitfrage dafür, das Szenario als Gedankenexperiment zu verwenden¹² und beispielsweise in verschiedenen Aspekten (z. B. Ausgangsproblem, KI-Potenziale, Auffassungen von Lehrpersonen etc.) zu variieren. In dieser Form lädt das Szenario dazu ein, darüber nachzudenken, ob und wenn ja, welche Gründe es für oder gegen die Umdeutung von Werten im Zuge der KI-Entwicklung und -Nutzung gäbe.

12 Vgl. Reinmann, Gabi, Gedankenexperimente als bildungstheoretisches Instrument in der Forschung zu Künstlicher Intelligenz im Hochschulkontext, in: Impact Free 56 (2014), zu finden unter https://gabi-reinmann.de/wp-content/uploads/2024/09/Impact_Free_58.pdf (abgerufen am 21. Jan, 2025), sowie Niesel, Dennis; Jelonnek, Stefan; Wilder, Nicolaus (in Druck), Gedankenexperimente als Methode pädagogischen Denkens – oder: Über die Notwendigkeit des Möglichen, Pädagogische Rundschau 2025.

Verändern sich im Kontext von KI-Technologien die Realisierungsbedingungen von Zielen und Werten der Hochschulbildung?

Was sagt unser Szenario dazu? In unserem Szenario werden zwei Arten von Bedingungen explizit angesprochen, die wichtig sind, um bisherige Ziele und Werte auch in Zeiten von KI zu realisieren: technische Infrastrukturen, die allen Studierenden, datenschutzrechtlich vertretbar, die Nutzung von KI ermöglichen, sowie (prüfungs-)rechtliche Bestimmungen, die Lehrpersonen flexible didaktische Entscheidungen und Lehrexperimente wie studentische Wahlfreiheit beim KI-Einsatz in Prüfungen ermöglichen. Implizit besteht eine weitere Realisierungsbedingung darin, dass Lehrpersonen über didaktische Kompetenz und Kreativität verfügen, denn: Am Ende sind alle Fragen zum Einsatz von KI in der Hochschulbildung auf der Mikroebene selbstredend didaktischer Art. Im Szenario werden in diesem Sinne die Bedingungen zur Realisierung von Zielen und Werten der Hochschulbildung teils vorausgesetzt, teils angepasst, damit diese auch mit KI erreichbar sind. Das Szenario zeigt ebenfalls, dass KI oftmals als **Treiber für Veränderungen** fungiert, die grundsätzlich sinnvoll wären, bisher aber vermutlich aus Mangel an „Leidensdruck“ nicht umgesetzt worden sind. Anders sähe die Einschätzung wohl aus, wenn wir ein Szenario hätten, in dem sich Ziele und Werte von Hochschulbildung verändern: Angenommen, wir würden unseren bisherigen Wert der Autor:innenschaft (weitgehend) aufgeben zugunsten einer hybriden Erstellung von Texten, bei der Mensch und Maschine ein nicht mehr zu entwirrendes Beziehungsgeflecht ergeben. In dem Fall würden Sarah und Yago in unserem Szenario andere (prüfungs-)didaktische Entscheidungen treffen bzw. andere studentische Kompetenzen fördern und in der Prüfungsgestaltung berücksichtigen. Sie würden einen neuen Wert etablieren – z. B. hybride Autor:innenschaft – und bräuchten dafür unterschiedliche Grade an **Änderungen technischer, rechtlicher und didaktischer Rahmenbedingungen**, die von den aktuellen deutlicher abweichen.

Wie ist die Leitfrage zu verstehen? Auch Leitfrage 2 lässt sich in zweifacher Weise lesen: als empirische Frage in dem Sinne, dass man via Befragungen oder Inhaltsanalysen prüfen könnte, wie die **derzeitigen Realisierungsbedingungen für den KI-Einsatz** an Hochschulen beschaffen sind und inwieweit sie sich in den letzten Jahren verändert haben, oder als philosophische Frage, ob und aus welchen Gründen wir diese Bedingungen in welcher Weise ändern sollten. Letzteres setzt allerdings voraus, dass wir schon zu dem Schluss gekommen sind, einzelne Ziele und Werte einer Veränderung unterziehen zu müssen. In beiden Fällen liegt unserer Einschätzung nach eine Herausforderung darin, dass „Realisierungsbedingung“ ein vieldeutiger Begriff ist. Wir haben diesen eingangs bereits in technische, rechtliche und didaktische Bedingungen aufgefächert. Nun kann man sich diese drei Aspekte auch als Achsen eines Koordinatensystems denken. Diese spannen einen Raum auf, in dem abzuwägen wäre, wo welche Vorgaben und Freiheitsgrade bei Technik und (Prüfungs-)Recht sowie bei der Entwicklung didaktischen Wissens und Könnens erforderlich oder begünstigend (oder hinderlich) sind, um bestimmte Ziele und Werte von Hochschulbildung zu erreichen.

Was folgern wir daraus? Auch in Bezug auf Leitfrage 2 tendieren wir dazu, aus dem Szenario ein Gedankenexperiment zu entwickeln, das in unterschiedlichen Variationen durchgespielt werden könnte. Dazu ließe sich das eben skizzierte Koordinatensystem mit den **drei Achsen Technik, Recht und Didaktik** heranziehen: Was wären die Folgen für Lehrende und

Studierende, wenn eine Hochschule auf diesen Achsen strenge Vorgaben machen und Freiheitsgrade einschränken würde? Welche Konsequenzen (z. B. für das Selbstverständnis der Lehrenden und Studierenden) sind denkbar, wenn die Vorgaben weit und die Freiheitsgrade groß wären? Was wäre, wenn es gar keine Vorgaben gäbe? Darüber hinaus ließen sich verschiedene Kombinationen von Vorgaben und Freiheitsgraden technischer, rechtlicher und didaktischer Art durchspielen.

Unsere eigene (gedanklich zu prüfende) These an dieser Stelle ist, dass der einzig sinnvolle Umgang mit steigender Komplexität und wachsender Dynamik im Zusammenhang mit KI sowohl für Lehrende als auch für Studierende darin zu suchen wäre, mehr Autonomie zuzulassen versus mehr Vorgaben zu machen.¹³ Das bedeutet im Gegenzug allerdings nicht, dass sich Hochschulleitungen zurückziehen sollten – im Gegenteil: Wenn es um die Realisierungsbedingungen im Zusammenhang mit disruptiven (technologischen) Innovationen wie KI geht, kommt Hochschulleitungen eine besondere Verantwortung zu, nämlich Orientierungsverantwortung. Das heißt: Hochschulleitungen müssten sich in Bezug auf Ziele und Werte positionieren und die **Gestaltung von Realisierungsbedingungen** darauf ausrichten: Will eine Hochschule der Integration von KI in sämtliche Lehr-, Forschungs- und Verwaltungsprozesse Priorität einräumen, erfordert das andere Realisierungsbedingungen als im Falle von Hochschulen, für die eine technologieunabhängige Autonomie Priorität hat. Für beide Positionen, die für Studierende und Lehrende und alle, die das werden wollen, transparent sein müssen, und das gesamte Kontinuum dazwischen lassen sich Argumente bzw. gute Gründe finden.

Hochschulleitungen obliegt also nicht nur die Aufgabe, technische Infrastrukturen zur Verfügung zu stellen, den rechtlichen Rahmen abzustecken und hochschuldidaktische Angebote zu machen, sondern innerhalb des sich dadurch ergebenden Möglichkeitsraums auch Orientierung anzubieten bzw. Positionierung anzuregen. Dabei kann es sich um Artefakte, Veranstaltungen, Vorbilder etc. handeln, die zum einen **nachvollziehbare Leitlinien** an die Hand geben und zum anderen Autonomie fördern, ohne dabei das Individuum zu überfordern oder orientierungslos zurückzulassen. Wird diese Orientierungsverantwortung wahrgenommen, ließe sich, so die optimistische Erwartung, eine gemeinsame Haltung und Identität entwickeln, die Ängsten und Überforderung ihrer Mitglieder entgegenwirken.

Wie können Ziele und Werte von Hochschulbildung durch verantwortungsvollen KI-Einsatz gestärkt werden?

Was sagt unser Szenario dazu? Unser Szenario ist bereits so verfasst, dass es zwei verschiedene Initiativen skizziert, in denen – die bisherigen Ziele und Werte von Hochschulprüfungen vorausgesetzt – ein verantwortungsvoller KI-Einsatz im Kontext von Hochschulprüfungen gestärkt wird: Die Wirtschaftsinformatikerin Sarah motiviert Studierende, mit KI zu arbeiten, aber in einer Form, dass sie das entstehende Produkt selbstständig (mündlich) vertreten können. Mit der Neugestaltung ihrer Prüfung stimmt sie diese mit veränderten Kompetenzanforderungen ab und löst gleichzeitig das beobachtete Pro-

13 Wilder, Nicolaus; Weßels, Doris, KI und das Ende des Einheitslehrplans? Eine kritische Analyse, in: Weiterbildung. Zeitschrift für Grundlagen, Praxis und Trends 35 (2025) Nr. 6.

blem eines **intransparenten und unkontrollierten KI-Einsatzes**. Der Politikwissenschaftler Yago ermutigt Studierende zu einer selbstbestimmten Entscheidung, was den KI-Einsatz angeht, bemüht sich dabei aber in beiden Fällen darum, dass Schreibkompetenz erworben wird, weil er diese in seinem Fach und für die Persönlichkeitsbildung für unentbehrlich hält. Mit der Möglichkeit, selbst zu entscheiden, ob KI genutzt wird oder nicht, gekoppelt mit der Bedingung, den **KI-Einsatz zu dokumentieren und zu reflektieren, fordert er Selbstverantwortung und Transparenz ein**. Auch er löst das Problem eines intransparenten und unkontrollierten KI-Einsatzes, tut dies aber auf anderem Wege als Sarah. Die derzeit bestehenden Ziele und Werte von Hochschulbildung werden mit beiden Maßnahmen gestärkt. Hätten wir ein Szenario, in dem sich Ziele und Werte von Hochschulbildung verändern – beispielsweise in Richtung hybride Autor:innenschaft –, müssten andere Maßnahmen bei der Prüfungsgestaltung ergriffen werden; parallel dazu wären rechtliche Vorgaben anzupassen und erforderliche technische Infrastrukturen sicherzustellen. In dem Fall wäre anzunehmen, dass sich die Prüfungskultur tiefgreifend verändert.

Wie ist die Leitfrage zu verstehen? Auch Leitfrage 3 ist aus unserer Sicht (im Nachhinein) dahingehend kritisch zu hinterfragen, ob man Ziele und Werte tatsächlich stärken kann wie z. B. Haltungen. Falls man das könnte, wäre auch hier zunächst zu klären, welche Ziele und Werte bei wem zu stärken sind und wie sich das realisieren ließe, ohne dogmatisch zu werden. Zu berücksichtigen wären also zunächst Antworten auf die erste Leitfrage, ob oder inwieweit Werte und Ziele der Hochschulbildung infolge von KI umgedeutet werden (sollen). Unser Szenario im Kontext Hochschulprüfungen enthält sich dieser kritischen Nachfrage und beugt sich gewissermaßen der gegenwärtig omnipräsenten Formulierung vom **„verantwortungsvollen Umgang mit KI“**. Mit dieser Formulierung wird genau genommen eine Prämisse gesetzt, die derzeit kaum in Zweifel gezogen wird, nämlich: KI ist die größte Chance und Gefahr für die Menschheit, also müssen wir verantwortungsvoll mit ihr umgehen, dann wird schon alles gut. So leicht einem dieses Heilsversprechen über die Lippen kommt, so schwierig ist dessen Konkretisierung: **Was genau heißt es, verantwortungsvoll mit KI umzugehen? Was sind die notwendigen Bedingungen für einen verantwortungsvollen Umgang mit KI? Ist Verantwortung lehrbar? Wer übernimmt sie?** Im Nachhinein würden wir die Frage, wie Ziele und Werte von Hochschulbildung durch verantwortungsvollen KI-Einsatz gestärkt werden können, so umformulieren: **Wenn man überzeugende Gründe für bestimmte Ziele und Werte von Hochschulbildung hat, was wäre zu tun, um sie trotz KI, mit KI oder wegen KI aufzubauen, konstruktiv weiterzuentwickeln und/oder zu bewahren?**

Was folgern wir daraus? Wir kommen auch hier wieder zu dem Schluss: Fruchtbare Diskussionen ließen sich vor allem mit einer gezielten **gedankenexperimentellen Variation des Szenarios** bezogen auf die letzte Leitfrage in Kombination mit den anderen beiden Leitfragen anregen. Je näher man im Szenario an bisherigen Zielen und Werten von Hochschulbildung bleibt, desto inkrementeller (vermutlich auch leichter umsetzbar) sollten die Folgerungen sowohl für die Gestaltung von Realisierungsbedingungen als auch für **Maßnahmen zur Stärkung eines verantwortungsvollen Umgangs mit KI** sein. Je mehr Werte bis hin zu Zielen von Hochschulbildung eine Transformation erfahren (Beispiel: hybride Autor:innenschaft), umso mehr werden wir vertraute Pfade verlassen und vermutlich auch Lehr-, Lern- und Prüfungskulturen verändern müssen – wenn denn der Wandel von Werten und Zielen erwünscht ist und begründet werden kann. Mit Blick auf die Frage nach der Verantwortung aber denken wir, dass eine weitere Leitfrage ergänzt werden müsste, denn: Wir gehen in der AG bislang

von der Frage aus, ob, und wenn ja, inwiefern wir unser Lehr- und Prüfungshandeln einschließlich der damit verbundenen Ziele und Werte durch die Existenz von KI verändern bzw. weiterentwickeln sollten. KI wird entsprechend als Ausgangspunkt gesetzt und es wird nach dem (verantwortungsvollen) Umgang mit den auf den Markt geworfenen KI-Systemen gefragt. Der Umgang mit KI vor dem Hintergrund von Zielen wie der Förderung einer kritisch-reflexive Haltung ist aber nur im unmittelbaren Lehr-Lerngeschehen realisierbar. So gesehen müsste die Lehre zum Ausgangspunkt gemacht und gefragt werden: Wie kann die Auseinandersetzung der Hochschulen mit KI in der Lehre zur werteorientierten Gestaltung von KI-Systemen beitragen?

Unser Szenario konzentriert sich weitgehend auf die Möglichkeit, dass KI die Gestaltung von Prüfungen und implizit von Lehre (als Konsequenz neuer Prüfungsformen) beeinflusst, indem die Technologie zum kritischen Hinterfragen bisheriger Routinen anregt. Nur indirekt deutet das Szenario die umgekehrte Chance an, dass die Auseinandersetzung der Hochschule mit KI beispielsweise zum Datenschutz (als einem Wert) dazu führt, dass zumindest die Implementierung von KI-Systemen an bisherigen akademischen Wertvorstellungen ausgerichtet wird (z.B. das datenschutzkonforme Aufsetzen von Large Language Models). Auszubauen wäre das Szenario jedoch, wenn wir genauer abzuwägen wollten, wie das Nachdenken über unsere Ziele und Werte in der Hochschulbildung den (dazu in der Regel kongruenten) gesellschaftlichen Werterahmen stabilisieren oder weiterentwickeln könnte, der gleichzeitig wieder der Rahmen für die Entwicklung und den Einsatz von KI-Systemen sein sollte.

Ausblick: Von der Integrität zur Verantwortung

Unser Szenario deutet an, wie stark KI bereits Einfluss auf konkrete Fragen der Prüfungsgestaltung und in der Folge auch der Lehrgestaltung nimmt. Dabei thematisiert das Szenario „nur“ die faktische oder unterstellte Delegation von prüfungsrelevanten Schreibaufgaben seitens der Studierenden auf generative KI, die mit dem (bisherigen) Wert der akademischen Integrität unvereinbar ist. Ebenso aber ließe sich gedanklich durchspielen, was passiert, wenn Lehrende beim didaktischen Handeln Aufgaben und damit auch Verantwortung an KI abgeben: sei es die Korrektur von Prüfungsleistungen, seien es die Zusammenstellung und Aufbereitung von Lern- und Prüfungsinhalten, sei es die Interaktion mit Studierenden im Prozess des Aneignens oder Einübens von Wissen und Können (via KI-gestützter tutorieller Systeme). Was in der Erwägung unseres Szenarios unter den drei Leitfragen der AG weitgehend außen vor blieb, sind in diesem Zusammenhang Fragen der **Rollen im Beziehungsgeflecht zwischen Lehrenden, Studierenden und KI** sowie die damit veränderten Anforderungen an die Übernahme von Verantwortung – ein Aspekt, der über unser Szenario hinausgeht und daher abschließend kurz erörtert werden soll.

Wenn wir etwa zunehmend KI-Systeme einsetzen, um Effizienzgewinne bei der Gestaltung und Umsetzung von Lehre zu erzielen, oder um Studierenden eine maßgeschneiderte Lernbegleitung für ihren Lernweg zu bieten, müssen wir uns fragen: Welche Wirkung hat das auf die sozialen Beziehungen sowie auf die Rollen von Lehrenden und Studierenden? Welche Rolle nimmt in diesem Beziehungsgeflecht die KI ein? Welche inhaltlichen und sozialen Anteile der Lehr-Lernbeziehung werden von KI-Systemen übernommen und welche von der Lehrperson und mit welcher Begründung? Selbstversuche in unserem Team legen nahe,

dass **aktuelle Chatbots in diesen Beziehungen „anthropomorphisiert“** werden und uns das Gefühl geben, in einem „echten“ sozialen Austausch zu sein. Wenn diese Erfahrungen flächendeckend gemacht werden: Welchen Folgen hat das? Wie weit ist der Weg bis zum menschenähnlichen Roboter, der den Anthropomorphismus konkret umsetzt? Welchen Einfluss hat das auf unsere Vorstellungen zum Verhältnis zwischen Mensch und Maschine und was bedeutet das für die hochschulische wie auch gesellschaftliche Wertebasis? Und schließlich: Wer ist in dieser Trias (Lehrende, Studierende, KI) wofür eigentlich (noch) verantwortlich?

Nun ist die Abgabe und/oder Diffusion von Verantwortung im Bildungskontext kein gänzlich neues, nur durch KI angestoßenes, Phänomen: Wenn etwa Lehrpersonen ihre Rolle vor allem als Coach und Lernbegleiter definieren, um damit das allseits wertgeschätzte Ziel der **Selbstverantwortung von Studierenden** zu stärken, besteht durchaus das Risiko, dass sie sich damit von der eigenen didaktischen Verantwortung „freikaufen“, indem sie diese (aus vermeintlich guten Gründen) an die Studierenden übertragen.¹⁴ Man mag einwenden, dass Studierende erwachsen sind und die Verantwortung ohnehin zu tragen haben, doch scheint diese Haltung zu einfach zu sein, denn sie lässt die Frage außen vor, welche Verantwortung Lehrende haben. **KI-gestützte Tutorielle Systeme ebenso wie Chatbots, die Lehrinhalte auf Knopfdruck generieren oder Qualifizierungsarbeiten begutachten, haben ein riesiges Potenzial, die Verantwortungsabgabe weiterzutreiben – ähnlich wie Studierende allzu schnell verleitet sein könnten, ihre Verantwortung etwa für schriftliche Artefakte an generative KI zu delegieren.**

Vor diesem Hintergrund sollten die Rollen mit ihren dazugehörigen Verantwortlichkeiten in der Hochschulbildung diskutiert und dann auch transparent definiert werden. Natürlich stehen auch hinter solchen Diskussionen und Entscheidungen jeweils Werte und Ziele, die auf diesem Wege explizit benannt werden müssten. Je nachdem, wie diese Rollen und Verantwortlichkeiten im Beziehungsgeflecht zwischen Lehrenden, Studierenden und KI aussehen, sind dann auch die dazu erforderlichen Bedingungen zu reflektieren. Einzubeziehen wären hier schließlich die **Anbieter von KI-Systemen**, wenn wir die obigen Ausführungen zur wertorientierten Gestaltung von KI-Systemen ernst nehmen.

Handlungsempfehlungen

Aus dem gedanklichen Durch- und Weiterspielen des Szenarios ließen sich zahlreiche Handlungsempfehlungen entwickeln – in Abhängigkeit vom inhaltlichen Fokus sowie der je eigenen Wertvorstellungen und Haltungen zum Thema Künstliche Intelligenz (KI). Die folgenden Empfehlungen sind ein stark kondensiertes Ergebnis von Reflexionen über das Szenario der Autor:innen. Sie basieren darauf, dass die eingangs erörterten akademischen Werte auch im Kontext von Prüfungsgeschehen (noch) gelten, also unter anderem: Integrität, Verantwortung, Transparenz, Autor:innenschaft, Selbstbestimmung, Partizipation, Chancengerechtigkeit, Offenheit bzw. Experimentierfreudigkeit.

(1) Hochschulleitungen übernehmen Orientierungsverantwortung: Sie positionieren sich

¹⁴ Ladenthin, Volker, Allgemeine Pädagogik, Baden-Baden 2022.

zum Thema KI und (a) erklären, welche Ziele und Werte auch unter „KI-Bedingungen“ für die Gestaltung von Prüfungen Bestand haben sollen, (b) richten technische, (prüfungs-) rechtliche und didaktische Bedingungen so aus, dass Prüfungen auch unter Nutzung von KI möglich werden, und (c) eröffnen Lehrenden und Studierenden Handlungs- und Entscheidungsspielräume im Kontext von Prüfungen, oder anders formuliert: Sie fördern die Handlungsautonomie, die Lehrende benötigen, um angemessen das Erreichen der jeweiligen Lernziele überprüfen zu können.

(2) Lehrende übernehmen Handlungsverantwortung: Sie schöpfen den ihnen gewährten Handlungsspielraum proaktiv aus und (a) berücksichtigen bei Entscheidungen in der Prüfungsgestaltung sowohl veränderte Kompetenzanforderungen von außen als auch fachwissenschaftliche Erfordernisse (infolge von KI), (b) geben Studierenden da, wo es möglich ist, Wahlfreiheit in Bezug auf die, vorab gemeinsam reflektierte, Nutzung von KI in Prüfungssituationen, (c) verzichten auf KI in Prüfungen dann, wenn gute Gründe dafür geltend gemacht werden können, und (d) erkennen und nutzen KI auch als Treiber ohnehin erforderlicher Veränderungen in Prüfungsroutinen und Prüfungskultur.

(3) Im Beziehungsgeflecht Studierende-Lehrende-KI geht Verantwortung nicht verloren: Lehrende und Studierende (a) reflektieren ihre Rollen im Beziehungsgeflecht mit KI, (b) klären, wer für was im Prüfungsgeschehen verantwortlich ist und (c) beziehen dabei den gesamten Prüfungszyklus ein: von der Gestaltung von Prüfungsaufgaben (Interaktion Lehrender-KI) über das Absolvieren der Prüfung (Interaktion Studierende-KI) bis zur Bewertung und Rückmeldung (Interaktion Lehrende-Studierender-KI), ohne die Maschine zu anthropomorphisieren.

(4) Meta-Empfehlung: Die extrem hohe Geschwindigkeit technologischer Entwicklungen im Bereich KI stellt insbesondere diejenigen vor große Herausforderungen, die eine Orientierung Verantwortung für andere und gleichzeitig den Anspruch haben, Zukunft aktiv zu gestalten und nicht bloß reaktiv Spielball fremder Interessen zu sein. Unsere tradierten Verfahren zum Umgang mit diesen Herausforderungen an Hochschulen scheinen hier an ihre Grenzen zu stoßen: Das Warten auf das Vorliegen empirischer Evidenzen etwa – das zumindest in den letzten Jahrzehnten maßgeblich für die Formulierung von Handlungsempfehlungen war –, dauert schlicht zu lange. So gibt es inzwischen zwar erste Ergebnisse etwa zu den Auswirkungen des Einsatzes von ChatGPT, doch zwischen diesen Aussagen und der Gegenwart liegen zwei Versionsnummern; Möglichkeiten und Grenzen des KI-Einsatzes haben sich seitdem stark verschoben. Die als wissenschaftlich geltenden Aussagen beispielsweise zu „Halluzinationen“ von Large Language Models (LLMs) sind so gesehen lediglich für die Vergangenheit evident, nicht jedoch für die Gegenwart. Damit rennen wir der technischen Entwicklung chronisch hinterher und geraten in einen potenziell handlungs lähmenden Zirkel: Wenn das ewige Vermessen dessen, was ist, als einziger Ansatz verfolgt wird, verunmöglicht dies eine Gestaltung der Zukunft. Eine Möglichkeit, aus diesem Zirkel ausbrechen und die Perspektive der wissenschaftlichen Reflexion ausgehend von gegebenen Evidenzen in die Zukunft zu richten, sehen wir – und das vorliegende Papier ist Ausdruck dessen – in dem gedankenexperimentellen Entwerfen und Durchspielen möglicher Szenarien. Auf dem Fundament empirischer Erkenntnisse (die Vergangenheit betreffend) lassen sich dann Projektionen der Zukunft erstellen, anhand derer auch und insbesondere Wertfragen und Zielvorstellungen – fundamental für Handlungsempfehlungen und der Empirie

per se nicht zugänglich – reflektiert, geprüft und gegebenenfalls neu ausgehandelt werden können. Neue Entwicklungen sind hier direkt integrierbar und lassen sich in Bezug auf Ihre Auswirkungen gedanklich abklopfen. Unsere methodische Empfehlung ist daher die Umkehrung oder drastischer Wiederbelebung der schon von Adorno 1964 beklagten „self-same[n] Schrumpfung des utopischen Denkens“, welches notwendig ist, um begründet, reflektiert und verantwortungsvoll Zukunft zu gestalten.

Gabi Reinmann, Nicolaus Wilder & Antje Michel

3.2 Szenario II: Autonomie

Ausgangslage

Das Team eines berufsbegleitenden, digitalen Zertifikats-Weiterbildungskurses an einer Hochschule möchte einen Chatbot und ein KI-basiertes Evaluationstool einsetzen, um Studierenden in Online-Seminaren schnell Feedback zu ihrem Fortschritt im Lernprozess zu geben und diesen zu bewerten. Gleichzeitig soll das System darstellen, wie effektiv und kosteneffizient die entsprechende Uni arbeitet. Seitens einzelner Lehrpersonen gibt es jedoch auch Vorbehalte gegenüber der Anwendung solcher Technologien, da diese mit erheblichen Risiken für die Studierenden verbunden sein könnten. Beispielsweise könnte das System potentiell zu **negativem Erwartungsdruck und Stresszuständen** bei Studierenden führen und es ergeben sich verschiedene Problematiken in Bezug auf den Schutz der Privatsphäre von Nutzenden.

Konstruktive Vision

Wenig später entscheidet die Hochschulleitung, dass sie in ein **KI-basiertes Feedback- und Evaluationstool** investieren möchte. Sie entscheidet sich für ein System namens Alvaluate, welches bereits an einigen Hochschulen in Europa und den USA im Einsatz ist. Das System verändert maßgeblich den Lernprozess von Studierenden und die Lehrpraxis von Lehrpersonen. Die Studierenden sind anfangs erfreut über die innovativen Funktionen des Tools, mit denen sich bisherige Lernerfolge nachzeichnen und offene Problemfelder identifizieren lassen. Lehrende wiederum haben nun einen viel detaillierteren Überblick über die Stärken und Schwächen der Studierenden und können diese Informationen nutzen, um den individuellen Lernprozess besser zu unterstützen.

Nach einiger Zeit jedoch verebbt die Euphorie gegenüber Alvaluate. Durch die häufig versandten (Push-)Nachrichten mit Hinweisen des System-Chatbots fühlen sich die Studierenden zunehmend gestresst. Auch wird ihnen durch den Wortlaut der Nachrichten immer klarer, wie allumfassend die Überwachung durch das System und damit auch durch die Lehrpersonen ist. Theoretisch lassen sich mit den Informationen, die den Lehrenden durch das System zur Verfügung gestellt werden, **Details über den Lebensalltag** der Studierenden ableiten. Es wird getrackt, wann und wie lange sie im System eingeloggt sind, wann sie Pause machen und von wo aus sie sich in den Webinaren zuschalten. In Prüfungen fällt zunehmend auf, dass sich einzelne Studierende in Bezug auf ihren Lernbedarf ausschließlich auf das nicht immer unfehlbare Empfehlungssystem von Alvaluate verlassen, anstatt

selbstkritisch und eigenständig Lücken in ihrem Lernstand zu identifizieren.

Vieles kommt durch ein **Datenleck** ans Licht, in welchem eine Studentin sich ins System gehackt, ein Profil als Lehrender erstellt und die dadurch verfügbaren Informationen einer Zeitung zugespielt hat. Parallel zu einer kontrovers geführten Debatte in der Presse formiert sich Widerstand unter Studierenden der Hochschule, in welchem Alvaluate in den Zertifikatsstudiengängen eingesetzt wurde. Die Studierenden fordern, den Vertrag mit Alvaluate umgehend zu kündigen und das System komplett aus dem universitären Alltag zu verbannen. Auch ein Teil der Lehrenden solidarisiert sich zunehmend mit den Studierenden. Um einem massiven Reputationsverlust und perspektivisch rechtlichen Klagen zu entgehen, geht die Hochschulleitung nach einer Weile auf die Forderung ein und löst die Geschäftsbeziehung mit Alvaluate auf. In Zukunft sollen technische Systeme durch ein gemeinschaftliches Gremium aus Studierenden-, Lehrenden- und Verwaltungsvertretung geprüft und gegebenenfalls zugelassen werden.

Ein Jahr später wird ein neues System eingeführt, Alssist, welches den Bedürfnissen der Betroffenen mehr entspricht und **ethische Erwägungen zum Schutz von Privatheit** und der Förderung der Autonomie von Studierenden berücksichtigt. Das neue System schließt Überwachungsfunktionalitäten dezidiert aus (privacy-by-design) und stellt eher die Kollaboration zwischen Studierenden sowie Feedback für die Lehrenden und die Hochschulleitung (Berechnung eines allgemeinen „Wohlfühlfaktors“ und anonymisierte Einschätzungen zur Qualität des Lehrangebots) in den Vordergrund. Auch hier wird ein **Chatbot zur Verfügung gestellt, der aber eher als Planungs- und Organisationsassistent funktioniert und ausschließlich auf expliziten Wunsch der Studierenden dezente Empfehlungen zum individuellen Kompetenzaufbau gibt**. Zudem lässt sich der Chatbot situativ ein- und ausschalten und die Nutzung des Systems ist ausschließlich freiwillig. In Bezug auf Lernstandsanzeigen der einzelnen Studierenden verzichtet das System auf vermeintlich detaillierte Angaben, da festgestellt wurde, dass KI-Systeme in diesem Bereich eine falsche Sicherheit vermitteln, sondern regt Studierende zu einer **selbstkritischen Reflexion des eigenen Lernstands an**. Da sich Alssist aber fast nahtlos in die Lehr- und Lernpraxis integriert und viele bisher umständliche Aufgaben vereinfacht, wird das System dennoch von fast allen Studierenden in den Zertifikatsstudiengängen genutzt und wenig später auf in allen anderen Lehrgängen angeboten und eingesetzt.

Diskussion

Das geschilderte Szenario kann im Hinblick auf mehrere Werte von Hochschulbildung diskutiert werden, zu denen im Bereich der Persönlichkeitsbildung natürlich Autonomie und Selbstbestimmung, im Bereich der (Fach-)Wissenschaft aber auch Mündigkeit und Kritikfähigkeit zählen.

Gehen wir davon aus, dass Hochschulbildung auf die Entwicklung autonomer, kritischer Denker:innen zielt, die ihre Lernprozesse selbst steuern können, bieten Assistenzsysteme wie Alvaluate und Alssist durch ihre Anleitungs- und Feedback-Funktionalität Lernenden eine **individualisierte Unterstützung** und erhöhen somit potenziell die Autonomie, etwa indem Studierende auf weniger Hilfe von außen angewiesen sind: Eine reflektierte Nutzung von Hinweisen, die durch automatisierte Lernstandsanalysen gegeben werden, kann Stu-

dierenden helfen, ein Gefühl für Lernbedarfe zu entwickeln – hier insbesondere im Hinweis auf bislang unbewusste Lücken im Lernprozess. Eine Gefahr besteht darin, dass sie, indem sie zu viele Vorgaben machen oder Studierende umfassend überwachen, die Autonomie wiederum einschränken. Teil eines erweiterten **Autonomie-Begriffs im KI-Zeitalter** könnte daher sein, die effektive Nutzung von Assistenzsystemen für sich selbst zu erlernen, ohne dabei die Selbststeuerung im Lernprozess zu verlernen. Ein System, das den Nutzenden ihre Lücken und Lernbedarfe vermeintlich en detail aufschlüsselt, birgt auf Seite der Nutzenden die Gefahr, dass Arbeits- und Lernbedarfe nicht mehr eigenständig erkannt werden können.

Datenschutz ist eine Voraussetzung für Autonomie. Die Einführung von Systemen wie Alvaluate zeigt, dass der **Datenschutz** insbesondere durch Überwachung und Datenspeicherung gefährdet ist. Zur Gewährleistung von Autonomie kann daher „privacy by design“ – also die Einbettung von Datenschutzprinzipien bereits in die technische Konzeption und Entwicklung von Systemen – einen wichtigen Beitrag leisten.

Bildung setzt voraus, dass eine Person selbst lernt; Lernprozesse individualisiert zu unterstützen, gilt daher als ein Ideal. KI-Systeme können hier einen Beitrag leisten, indem sie individuell auf Lernstände und -bedarfe reagieren und Feedback geben. Für die Entwicklung von **Mündigkeit und Kritikfähigkeit** ist es aber essenziell, dass Studierende die Fähigkeit erlernen, selbstständig die nächsten Schritte in ihrem Lernprozess zu erkennen. KI-Systeme sollten darin unterstützen, nicht jedoch die vollständige Steuerung der Lernschritte übernehmen und Studierende in ihrer Verantwortung für das eigene Lernen und somit in ihrer Bildung zu beschränken.

Weitere Werte, die im Szenario tangiert werden, sind Transparenz und Fairness: Die Algorithmen, die hinter KI-basierten Tools im Hochschuleinsatz stehen, müssen transparent und fair sein. Die mangelnde Transparenz und das Datenleck im Szenario Alvaluate zeigen, dass hier erhebliche Gefahren bestehen. Hochschulen sollten sicherstellen, dass ihre KI-Systeme auf ethischen Prinzipien basieren, um vertrauenswürdige KI-Nutzungsszenarien zu gestalten.

Handlungsempfehlungen

Aus dem dargestellten Szenario ergeben sich für den verantwortungsvollen und konstruktiven Umgang mit KI-basierten **Feedback- und Evaluationstools** an Hochschulen folgende Handlungsempfehlungen:

[1] Hochschulleitungen sollten klare Standards und Richtlinien zur Nutzung KI-basierter Systeme definieren und dabei insbesondere (a) **Datenschutz und Privatsphäre** gewährleisten, etwa durch datenschutzkonforme und transparente Systemauswahl („privacy by design“), (b) die Nutzung dieser Systeme von der Zustimmung der Studierenden abhängig machen und (c) die Einführung und Nutzung solcher Systeme kritisch durch Gremien begleiten und evaluieren lassen. Außerdem sollten sie (d) die Lehrenden durch gezielte Maßnahmen, Angebote und Fortbildungen zum kompetenten und kritischen Umgang mit KI-basierten Systemen befähigen und deren Bereitschaft zum Aufbau dieser Kompetenzen aktiv fördern.

(2) Lehrende übernehmen eine aktive Rolle bei der Integration KI-basierter Tools in ihre didaktische Praxis, indem sie: (a) gemeinsam mit den Studierenden klar festlegen, welche Informationen durch KI erfasst werden dürfen, (b) transparent kommunizieren, wie die erhobenen Daten genutzt und interpretiert werden und welche Rückschlüsse daraus auf den Lernprozess gezogen werden können oder eben nicht, und (c) KI-gestütztes Feedback stets ergänzend und unterstützend einsetzen, jedoch niemals die Verantwortung für die Beurteilung von Lernständen oder -erfolgen an die KI abgeben.

(3) Der Schutz der Studierenden vor unerwünschter Überwachung und stressauslösenden Kontrollmechanismen muss oberste Priorität haben. Hierfür ist es erforderlich, (a) den Umfang der Datenerhebung durch KI-Systeme strikt auf ein notwendiges Minimum zu begrenzen, (b) Daten stets anonymisiert zu verarbeiten, um Rückschlüsse auf individuelles Verhalten auszuschließen, und (c) transparent gegenüber den Studierenden offenzulegen, welche Daten zu welchem Zweck und für welchen Zeitraum erhoben und verarbeitet werden.

(4) Die Autonomie und Mündigkeit der Studierenden sollen durch den Einsatz KI-basierter Systeme gestärkt, nicht geschwächt werden. Dies gelingt durch Maßnahmen, wie: (a) die Vermeidung vermeintlich allwissender, deterministischer Aussagen zum Lernstand und stattdessen Anregungen zu selbstkritischer Reflexion der eigenen Lernprozesse, (b) Förderung einer bewussten, selbstbestimmten Nutzung der Systeme, bei denen Studierende ihre eigenen Ziele setzen und verfolgen können, und (c) klare Kommunikation darüber, dass die durch KI-Systeme gelieferten Einschätzungen keine unumstößlichen Wahrheiten darstellen, sondern Impulse für eigenständige Reflexion und selbstbestimmtes Lernen bieten sollen.

(5) Hochschulleitungen und Lehrende sollten die Erfahrung aus der Nutzung von KI-Systemen aktiv nutzen, um institutionelle Lernprozesse anzustoßen: Dies umfasst, (a) regelmäßig systematische Reflexionen und Analysen bisheriger KI-Projekte vorzunehmen, (b) Erkenntnisse in künftige Entscheidungen über KI-Einsatz und Leitlinien einfließen zu lassen und (c) eine offene Fehlerkultur zu etablieren, die Innovation und konstruktiv-kritische Reflexion im Umgang mit KI-Technologien fördert.

Ulrike Tippe & Martin Wan

3.3 Szenario III: Forschung

Ausgangslage

Dr. Jona Müller leitet an einem Institut für angewandte Informatik eine Forschungsgruppe zu Technologien des Natural Language Processings (NLP). Ihre Gruppe entwickelt unter anderem KI-Algorithmen, die dabei helfen sollen, **Hassrede und den Aufruf zu Gewalt in Online-Kommunikationsdaten** zu erkennen.

Jona hat in der Vergangenheit verschiedentlich mit großen Verlagen und Internet-Konzernen gearbeitet, die solche Algorithmen für das Filtern von Inhalten auf (sozialen) Medienplattformen nutzen. Seit einem Jahr kooperiert sie zudem das erste Mal mit deutschen Polizeibehörden. Technisch gesehen sind die zu entwickelnden Verfahren ähnlich wie jene, an

denen die Forschungsgruppe vorher gearbeitet hat, aber durch die neuen Partner aus der Polizei verändern sich die Rahmenbedingungen für die Forschung in erheblichem Maße. Vor einiger Zeit erschien ein Artikel in einer angesehenen Wochenzeitung, der sich kritisch mit Überwachungstechnologien im KI-Zeitalter auseinandersetzte. Ihr Projekt wurde namentlich als Beispiel für **problematische Kooperationen zwischen Forschenden und Strafverfolgungsbehörden** erwähnt. Seither stehen ihr einige ihrer Kolleg:innen aus dem Institut kritisch gegenüber. Ein wichtiger Projektmitarbeiter hat mit dem Verweis gekündigt, er könne eine Weiterarbeit nicht mit seinem Gewissen vereinbaren. Es war bisher unmöglich, eine neue kompetente Mitarbeiterin für das Projekt zu finden und sie hat das Gefühl, dass dies neben dem Fachkräftemangel auch mit dem Anwendungskontext und schlechten Image des Vorhabens zu tun haben könnte.

Jona ist sich grundsätzlich der Verantwortung bewusst, die mit solchen Technologien einhergeht. Auch sie beginnt, die ethischen Implikationen ihrer Forschung systematisch zu hinterfragen. Ihr ist beispielsweise klar, dass eine unreflektierte Einführung der KI potentiell zu **Menschenrechtsverletzungen** führen könnte, etwa durch Datensätze, die einen Bias haben und so zu diskriminierenden Ergebnissen führen. Zudem gibt es keine klaren gesetzlichen Richtlinien, die den Einsatz solcher Technologien in der Strafverfolgung regeln. Jona fühlt sich zunehmend überfordert und weiß nicht so recht, wie sie den an sie herangetragenen Kritikpunkten außerhalb der Technikentwicklung souverän begegnen soll. Das Projekt endet ein Jahr später mit mittelmäßigen Ergebnissen und einem unguten Gefühl im Magen.

Konstruktive Vision

So sucht Jona aktiv nach Wegen, um ihre offenen Fragen besser zu verstehen. Sie engagiert sich vermehrt in Workshops und Netzwerken, in welchen **Fragen der algorithmischen Fairness**, der Transparenz und der ethisch-rechtlichen Rahmenbedingungen von KI und Überwachungstechnik diskutiert werden. Sie kommt ins Gespräch mit Ethiker:innen, Sozialwissenschaftler:innen, Jurist:innen, aber auch Vertreter:innen von Menschenrechtsorganisationen und der Industrie.

Gemeinsam mit solchen Akteuren bewirbt sie sich noch einmal erfolgreich für Fördermittel, diesmal zur Entwicklung von **Technologien zur Prävention und Verfolgung von Cyberkriminalität**. Das Projekt bekommt von der Ethik-Kommission der Uni ein Votum, wie es heute etwa bei Medizinprojekten Standard ist.

Im Gegensatz zum vorigen Projekt wird nun von Anfang an an einem Setting integrativer Forschung gearbeitet. Das bedeutet, dass Ethiker:innen, Sozialwissenschaftler:innen, Jurist:innen, **Vertreter:innen von Menschenrechtsorganisationen** und der Industrie gemeinsam an der Entwicklung der Technologie mitarbeiten und dass technische Komponenten iterativ unter Berücksichtigung ethischer, sozialer und rechtlicher Fragen entwickelt werden.

In einem der Workshops des Projektkonsortiums unterhält sich Jona mit Dr. Fiona Schmidt, einer Medienwissenschaftlerin. Fiona forscht zur gesellschaftlichen Akzeptanz digitaler Technologien und zur Wissenschaftskommunikation. Sie schlägt vor, gemeinsam ein Internet-Portal zu entwickeln, welches problematische Seiten von KI-Technologie im Überwachungskontext diskutiert und Möglichkeiten zu Bürger:innen-Feedback und Dialog bereit-

stellt. Es gelingt den beiden, Geld für die Entwicklung der Plattform aufzutreiben und das Wissenschaftskommunikationsprojekt läuft nun parallel zum Vorhaben der Technologieentwicklung. Das Internet-Portal entwickelt sich zu einer zentralen Anlaufstelle für Information und **Diskussion von Perspektiven digitaler Überwachung**. Die kollaborativ zustande gekommenen Inhalte und Online-Diskussionen informieren die Entwicklung der KI-Technologie und begleitende Maßnahmen wie die Konzeption geeigneter Governance-Strukturen, welche eine ethisch-vertretbare und verantwortungsvolle Nutzung der Technologie unterstützen. Das breit abgestützte Vorgehen der KI-Technikgestaltung kann zwar nicht alle möglichen Schattenseiten digitaler Überwachung abfedern. Aber für Jona steht dies auch nicht im Vordergrund.

Sie schätzt insbesondere, dass die Interaktionen mit interdisziplinären Partnern und der interessierten Öffentlichkeit neue Fragestellungen an sie herantragen, die ihre Forschung inspiriert und sie mehr mit der Praxisebene der KI-Nutzung verknüpft. Zudem ist sie glücklich darüber, dass sie ihre Arbeit wieder besser mit ihrem Gewissen vereinbaren und womöglich positive Impulse für **verantwortungsvollere Technologiegestaltung** in die wissenschaftliche und öffentliche Debatte einbringen kann.

Handlungsempfehlungen

In sensiblen Forschungs- und Entwicklungsbereichen sind interdisziplinäre Teams, systematische Risiko-Analysen, die Einbindung relevanter Stakeholder und aktive Wissenschaftskommunikation wichtig. Diese Maßnahmen tragen dazu bei, das Vertrauen verschiedener gesellschaftlicher Gruppen in die Integrität der Forschung und die verantwortungsvolle Entwicklung sowie den Einsatz von KI-Technologien zu stärken. Gleichzeitig fördern sie die gesellschaftliche Akzeptanz der Technologien, indem sie Transparenz schaffen und Partizipationsmöglichkeiten eröffnen.

(1) Ethische Implikationen konkret, kontextabhängig und frühzeitig beurteilen („Ethics by Design“): Als Basistechnologie kann KI sehr vielseitig eingesetzt werden und daraus ergeben sich zahllose Gestaltungsmöglichkeiten für die Systeme, die Mensch-Maschine-Interaktion und den Nutzungskontext, die darüber entscheiden, ob es zu einer wünschenswerten oder problematischen Nutzung kommt. Dies schließt unter anderem Themen wie den Schutz von Privatheit, algorithmische Fairness, Transparenz und Zugang, menschliche Kontrolle, Risiken von Dual Use und Manipulation sowie (ökologische, soziale und ökonomische) Nachhaltigkeit mit ein. In den einzelnen Fachdisziplinen sollten ethische Implikationen frühzeitig und möglichst bezogen auf konkrete Anwendungsfelder (idealerweise auf eine intendierte Anwendung) beurteilt werden, um Fehlanwendungen zu minimieren und einen verantwortungsvollen Einsatz zu fördern.

(2) Inter- und transdisziplinäre Zusammenarbeit fördern: Die Herausforderungen und neuen Aufgaben, die mit Technologien einhergehen, können von keiner Disziplin allein angegangen werden, sondern betreffen alle wissenschaftlichen Disziplinen und brauchen interdisziplinäre Zusammenarbeit.

Zugleich werden gerade bei Technik außerwissenschaftliche und verschiedene gesellschaftliche Akteure besonders relevant, d. h. die transdisziplinäre Zusammenarbeit rückt hier in

den Fokus. Das betrifft die Zusammenarbeit z. B. mit Kunst und Design, Akteuren aus dem Kulturbereich, religiösen Gemeinschaften, Politik, NGOs, Vereinen und Stiftungen.

Die technologischen Veränderungen betreffen alle Lebensbereiche, transformieren Gesellschaft und stellen so eine gesamtgesellschaftliche Gestaltungsaufgabe dar; umgekehrt benötigt die Wissenschaft normative Vorgaben aus der Gesellschaft etwa über gewünschte Anwendungsszenarien oder Beschaffenheit von Trainingsdaten (Stichwort: Bias). Beim transdisziplinären Arbeiten werden Forschungsgegenstand und Forschungsfragen von Wissenschaft und Gesellschaft gemeinsam entworfen und die Wissenschaft lenkt ihren Blick auf aktuelle Herausforderungen in der Gesellschaft. Wenn Wissenschaft und Gesellschaft gemeinsam lernen und Zukunftsvorschläge entwerfen, verändern sich Forschungsmethoden und Kommunikation. Es ist daher wichtig, neue partizipative Formate zu erproben und eine gemeinsame Sprache/Verständigungsbasis zu finden/erlernen (also Konzepte, Begriffe etc. klären...).

[3] Risiko-Analyse und Stakeholder-Einbindung in sensiblen Feldern: Die Einbindung von Stakeholdern in Forschungs- und Entwicklungsprozesse bietet Chancen und Herausforderungen. Externe Akteure wie staatliche Institutionen, Unternehmen oder Akteure der Zivilgesellschaft können praxisnahe Perspektiven und Ressourcen bereitstellen, was die gesellschaftliche Akzeptanz und ethische Akzeptabilität von Technologiegestaltung und -anwendung fördern kann. Gleichzeitig entsteht ein Spannungsfeld, da die Einbindung von Stakeholdern zu problematischer Beeinflussung oder auch zu finanziellen und infrastrukturellen Abhängigkeiten führen kann. Hier müssen gesellschaftliche Teilhabe und Wissenschaftsfreiheit gegeneinander abgewogen werden.

Stakeholder verändern die Forschung, indem sie Themen und Methoden mitprägen. Kooperationen mit großen Technologieunternehmen können innovative Ansätze fördern, bergen aber das Risiko, die Forschungsagenda an externe Interessen anzupassen. Ähnliche Dynamiken entstehen bei staatlichen Partnern, deren politische Ziele die Forschungsprioritäten beeinflussen können.

Prozedurale Ansätze wie ethics-by-design¹⁵, value-sensitive design und integrierte Forschung können dazu beitragen, die Vorteile solcher Kooperationen zu nutzen und das Risiko unerwünschter Entwicklungen zu minimieren.

[4] Wissenschaftskommunikation – Wissen verantwortungsvoll und adäquat an Betroffene (gesellschaftliche & politische Akteure, ggfs. auch breite Öffentlichkeit) kommunizieren und Partizipation ermöglichen: Die rasanten technologischen Entwicklungen führen zu gesellschaftlicher Unsicherheit, Orientierungsbedarf und Polarisierung zwischen Technikangst und Technikeuphorie. Diese Verunsicherung stellt Individuen, Institutionen und die Demo-

¹⁵ Ethics-by-design bedeutet, dass ethische Überlegungen von Beginn an systematisch in den Entwicklungsprozess von Technologien integriert werden, statt sie erst nachträglich zu berücksichtigen. Value-sensitive design verfolgt einen ähnlichen Ansatz, indem es gesellschaftliche Werte – wie Fairness, Transparenz oder Datenschutz – bereits in die Gestaltung technischer Systeme einfließen lässt, sodass diese nicht nur funktional, sondern auch wertorientiert entwickelt werden.

kratie vor Herausforderungen, insbesondere wenn Menschen ihre Identität, Handlungsfähigkeit oder berufliche Zukunft bedroht sehen. Medien verstärken diese Spannungen teilweise durch zugespitzte Berichterstattung, weshalb wissenschaftlich fundierte und differenzierte Medienarbeit und Wissenschaftskommunikation unerlässlich sind.

Die Forschung trägt eine besondere Verantwortung, Wissen nicht nur zu generieren, sondern auch angemessen und pluralistisch zu kommunizieren. Es gilt, technologische Deutungsmacht nicht allein großen Unternehmen oder einzelnen Akteuren zu überlassen, sondern Gesprächsangebote zu schaffen, die Partizipation ermöglichen. Partizipative Schnittstellen („participatory interfaces“) können der Gesellschaft Feedback- und Mitgestaltungsmöglichkeiten bieten, wodurch Transparenz und kritisches Mitdenken gefördert werden. Dabei ist wichtig zu betonen, dass Aufgaben der Wissenschaftskommunikation nicht allein durch die Forschenden gestemmt werden sollen, sondern dass unterstützende Strukturen für Dialog und Wissenstransfer zur Verfügung stehen.

Simon Hirsbrunner, Aljoscha Burchardt & Anna Puzio

4 Handlungsperspektiven für die Hochschulen im Kontext von KI: Möglichkeiten der Kompetenzentwicklung

Die übergeordnete Frage, die allen geschilderten Szenarien zugrunde liegt, lautet: Wie sollen sich Hochschulen im Kontext von Künstlicher Intelligenz entwickeln? Auf welche Weise müssen Kompetenzen der Hochschulangehörigen herausgebildet werden, damit die Ziele **(Fach-)Wissenschaft, Persönlichkeitsbildung und Arbeitsmarktvorbereitung** sowie die damit verbundenen Werte gestärkt werden? Wer sind Handlungsakteure dafür, wie kann deren Verantwortungsbereich beschrieben werden und welche Ressourcen sind dafür nötig?

Antworten auf diese Fragen legen bereits die Szenarien nahe, sie sollen in diesem Abschnitt aber auf einer allgemeineren Ebene expliziert werden. Das soll Hochschulen **Hilfestellung für ihre Entwicklung** geben, die angesichts von KI-Technologien notwendig wird. Die Gestaltung von Studium und Lehre zur Kompetenzbildung sollte sich primär an Zielen und Werten von Hochschule orientieren und dabei technologische Entwicklungen reflektiert einbeziehen.

4.1 Übergeordnete Handlungsempfehlungen

KI als strategisches Thema identifizieren

Damit Hochschulen den tiefgreifenden Wandel durch Künstliche Intelligenz aktiv mitgestalten können, bedarf es einer gezielten strategischen Herangehensweise. Ein zentraler Schritt ist die **bewusste Identifikation von KI als Entwicklungsthema auf institutioneller Ebene**. Hochschulen sollten durch eine strukturierte Herangehensweise sicherstellen, dass Lehrende, Forschende und Studierende gleichermaßen die Möglichkeit haben, sich mit KI auseinanderzusetzen und befähigt werden, sie reflektiert einzusetzen.

Ziel muss es sein, das Bewusstsein für das Potenzial und die Herausforderungen von KI zu fördern. Hierzu können hochschulweite Sensibilisierungs- und Schulungsformate beitragen. Darüber hinaus sollten **Fachbereiche und Organisationseinheiten** dazu ermutigt werden, eigene Leitlinien für den Umgang mit KI zu entwickeln, die (fach-)spezifische Herausforderungen und Ziele adressieren. Wichtig ist, dass Hochschulen hierbei einen offenen, explorativen Ansatz wählen, der **Raum für unterschiedliche Integrationstiefen** lässt – von der Schaffung KI-freier Räume bis hin zur umfassenden Implementierung von KI in akademische Prozesse. Entscheidend ist, dass dieser Prozess hochschulautonom und inhaltlich fundiert gestaltet wird, anstatt vorschnell von externen politischen und finanziellen Rahmenbedingungen determiniert zu werden.

Eine erfolgreiche Implementierung solcher Konzepte kann dadurch begünstigt werden, dass Hochschulen **Räume für interdisziplinären Austausch** schaffen, die sich außerhalb der gewohnten Gremien- und Verwaltungsstrukturen bewegen. Formate wie hochschulweite Diskussionsforen oder themenbezogene Innovationsveranstaltungen könnten hierfür ein probates Mittel sein.

Erfolgreiche Hochschulstrategien im Umgang mit KI sollten in einem **iterativen Prozess weiterentwickelt und regelmäßig überprüft** werden. Durch einen bewussten Fokus auf hochschulspezifische Bedarfe, eine offene Gestaltungsmethodik und die Förderung einer diskursiven Kultur lässt sich eine nachhaltige und reflektierte Integration von KI an Hochschulen realisieren. Entscheidend ist dabei die Stärkung der **Selbstbestimmung aller Hochschulangehörigen** im Umgang mit KI, um nicht nur Potenziale optimal zu nutzen, sondern auch Risiken angemessen zu begegnen.

Allgemeines Verständnis von KI fördern

Alle Hochschulangehörigen sollten ein **allgemeines Verständnis von Künstlicher Intelligenz entwickeln**, um einen verantwortlichen und didaktisch sinnvollen Einsatz von KI-Werkzeugen zu gewährleisten. Dabei sollte der Fokus nicht nur auf der praktischen Nutzung von KI, sondern im Wesentlichen **auf der kritischen Reflexion der zugrundeliegenden Technologie liegen**. Hier ist insbesondere das **Verständnis von generativer KI als statistischer Software** gemeint: Sprachmodelle – einschließlich sogenannter „Reasoning“- oder „Deep Research“-Modelle – basieren nach wie vor primär auf statistischen Verfahren der Mustererkennung: Sie generieren Texte, indem sie Wortfolgewahrscheinlichkeiten aus ihren Trainingsdaten und gegebenenfalls aus zusätzlichen externen Quellen zur Laufzeit („Retrieval-Augmented Generation“, RAG) berechnen, im Fall von fortgeschrittenen Modellen kombiniert mit Ansätzen symbolischer, also regelbasierter KI.¹⁶

Das ist wichtig, weil bei modernen KI-Modellen der schon von Joseph Weizenbaum beschriebene ELIZA-Effekt,¹⁷ also die Gefahr der Selbsttäuschung aufgrund der Illusion eines verstehenden Gegenübers, besonders hoch ist. **Weil der generierte Output in den meisten Fällen semantisch und inhaltlich sinnvoll erscheint, ist man geneigt, dem System zu vertrauen und ihm menschenähnliches Verstehen zu unterstellen.** Gleichwohl sind und bleiben KI-Systeme Software, die nach dem Prinzip Input-Verarbeitung-Output funktionieren. Dass die Verarbeitung von KI-Systemen aufgrund der schieren Masse von Gewichten und Parametern hochkomplex und nicht mehr manuell nachvollziehbar ist, bedeutet nicht, dass das, was in KI-Systemen passiert, prinzipiell unerklärlich ist. Die mangelnde Transparenz aufgrund der schieren Größe und der damit einhergehenden Undurchschaubarkeit der Gewichte von KI-Modellen muss aber Teil des Problembewusstseins ihrer Anwender sein.

Damit einher geht auch ein Bewusstsein für die Fehlerquellen von KI-generierten Outputs. Als statistische Modelle „verstehen“ KI-Systeme keine Bedeutungen und Sinnzusammenhänge, sondern replizieren Inhalte allein aufgrund statistischer Wahrscheinlichkeiten. Anstelle von Zusammenfassungen generieren Sprachmodelle deshalb auch oftmals lediglich

16 Vgl. Se, Ksenia, How Do Agents Plan and Reason?, zu finden unter: <https://huggingface.co/blog/Kseniase/reasonplan> (abgerufen am 11. Mrz. 2025).

17 Vgl. Wan, Martin, Kann KI bewusst sein? Anmerkungen zur Diskussion um LaMDA (Juli 2022), zu finden unter: <https://digiethics.org/2022/07/22/kann-ki-bewusst-sein-anmerkungen-zur-diskussion-um-lamda/> (abgerufen am 27. Feb. 2025).

gekürzte Texte, die wesentliche Informationen auslassen.¹⁸ Für das wissenschaftliche Arbeiten aber ist ein Verständnis von Sinnzusammenhängen über verschiedene Quellen hinweg sowie **Quellenkritik** essenziell. Während Sprachmodelle teils richtig aussehende, aber nicht existente Quellen erfinden,¹⁹ scheitern aktuelle „Reasoning“- und „Deep Research“-Modelle vielfach an der Unterscheidung seriöser und unseriöser Quellen oder belegen mit der angegebenen Quelle nicht, was sie zu belegen behaupten.²⁰ Auch bei erwartbaren künftigen Verbesserungen der Modelle, etwa durch Zugriff auf mehr und seriösere Quellen, bleibt eine Kritikfähigkeit gegenüber den Outputs und den angegebenen Quellen unverändert wichtig.

Das Verständnis von KI-Modellen als statistischer Software ist somit eine Voraussetzung für die Beurteilung nicht nur von konkreten KI-Outputs, sondern weiter gefasst von sinnvollen und sinnlosen bzw. potenziell gefährlichen Einsatzgebieten von KI. Wer versteht, dass KI nicht „versteht“, kann die Gefahren riskanter Anwendungen von KI, etwa in vollautomatisierten Entscheidungsprozessen, besser erkennen. Hilfreich hierfür kann auch eine **KI-Erkundungskompetenz** sein, also das Ausprobieren von KI-Modellen und die Fähigkeit, diese – etwa durch geschicktes Prompting – an ihre Grenzen zu bringen. Ein regelmäßiger Umgang mit den einschlägigen Modellen kann auch helfen, eine kreative KI-Kompetenz, also den Einsatz von KI zur Unterstützung menschlicher Kreativität, zu schulen. Wünschenswert wäre hierfür ein flächendeckender, zuverlässiger und datensparsamer Zugang zu gängigen KI-Modellen an deutschen Hochschulen.

Quellen- und Ideologiekritik üben

Die klassische akademische Kompetenz der Quellenkritik gewinnt durch KI nochmals an Bedeutung. Insbesondere kommerzielle KI-Anbieter legen im Hinblick auf ihre Trainingsdaten wenig Transparenz an den Tag. Dieser Transparenzmangel setzt sich bei durch KI-Modelle generierten Inhalten fort. Zur **Quellenkritik** gehört in diesem Zusammenhang auch eine Sensibilität für das Modelltraining mittels urheberrechtlich geschützter Inhalte und deren anschließende (teilweise) Reproduktion. „Reasoning“- und „Deep Research“-Modelle ent-

18 Vgl. Wierda, Gerben, When ChatGPT summarises, it actually does nothing of the kind, zu finden unter: <https://ea.rna.nl/2024/05/27/when-chatgpt-summarises-it-actually-does-nothing-of-the-kind/> (abgerufen am 28. Feb. 2025).

19 Vgl. Wan, Martin, Warum ChatGPT nicht das Ende des akademischen Schreibens bedeutet, zu finden unter: <https://digiethics.org/2023/01/03/warum-chatgpt-nicht-das-ende-des-akademischen-schreibens-be-deutet/> (abgerufen am 28. Feb. 2025) sowie Ders., Warum ChatGPT (auch weiterhin) nicht das Ende des akademischen Schreibens bedeutet, zu finden unter: <https://digiethics.org/2024/01/23/warum-chatgpt-auch-weiterhin-nicht-das-ende-des-akademischen-schreibens-bedeutet/> (abgerufen am 28. Feb. 2025).

20 Vgl. Jones, Nicola, OpenAI's 'deep research' tool: is it useful for scientists? Zu finden unter: <https://www.nature.com/articles/d41586-025-00377-9> (abgerufen am 28. Feb. 2025) sowie Major, Eddie, RAGs, Reasoning and Deep Research: What's new in AI and what might it mean for teaching in 2025? Zu finden unter: <https://www.adelaide.edu.au/learning/news/list/2025/02/21/rags-reasoning-and-deep-research-whats-new-in-ai-and-what-might-it-mean-for> (abgerufen am 28. Feb. 2025).

scheiden automatisiert, welche Quellen sie für die Generierung ihrer Texte hinzuziehen, so dass die Benutzenden eine Urteilsfähigkeit im Hinblick auf deren Qualität besitzen müssen.

Eine **Awareness gegenüber möglichen Biases in den Trainingsdaten setzt auch die Fähigkeit zur Ideologiekritik voraus**: Je nach Anbieter, Trainingskorpus und Feintuning weisen Modelle unterschiedliche Biases und Verzerrungen auf.²¹ Zudem sollte ein Bewusstsein gegenüber sozial und ökologisch problematischen Aspekten großer KI-Modelle (Ausbeutung von Clickworkern im globalen Süden, Ressourcenverbrauch der notwendigen Hardware) vorhanden sein, um KI-Einsatz unter ethischen Gesichtspunkten zu bewerten.

In den Bereich der Ideologiekritik fällt auch die **Sensibilität gegenüber PR von KI-Unternehmen vs. tatsächlichen Fortschritten in der KI-Forschung**. Marketingbegriffe wie AGI („Artificial General Intelligence“) oder „Emergenz“ befördern die Mystifizierung von KI und die Wahrnehmung von KI-Modellen als Wissens- und Entscheidungsmaschinen mit menschenähnlichen Fähigkeiten.²² Gleichzeitig konnte bislang weder überzeugend gezeigt werden, dass KI mehr sei als ausgeklügelte Software²³, noch welche Bedingungen erfüllt sein müssten, damit Software menschenähnliche Fähigkeiten entwickelt: Es gibt keinen philosophischen oder wissenschaftlichen Konsens über die genaue Funktionsweise des menschlichen Gehirns²⁴, über die Frage nach dem Bewusstsein oder die Beantwortung des

21 Das chinesische Modell DeepSeek, auch in unzensurierter Version, antwortet aus einer der chinesischen Staatsführung genehmen Perspektive, vgl. Lee Myers, Steven, DeepSeek's Answers Include Chinese Propaganda, Researchers Say, in: New York Times (online), zu finden unter: <https://www.nytimes.com/2025/01/31/technology/deepseek-chinese-propaganda.html> (abgerufen am 28. Feb. 2025). Elon Musks Grok 3 verrät sich bei der Frage nach den größten Verbreitern von Fake News, die Antwort sei schwierig, weil das Modell angewiesen sei, Musk und Trump bei der Antwort herauszuhalten, vgl. Hanfeld, Michael, Die KI von Elon Musk verplappert sich, in: FAZ (online), zu finden unter: <https://www.faz.net/aktuell/feuilleton/grok-3-laesst-musk-und-trump-bei-frage-nach-fake-news-weg-110319792.html> (abgerufen am 28. Feb. 2025).

22 Vgl. auch Weizenbaum, Joseph: „Der meiste Schaden, den der Computer potentiell zur Folge haben könnte, hängt weniger davon ab, was der Computer tatsächlich machen kann oder nicht kann, als vielmehr von den Eigenschaften, die das Publikum dem Computer zuschreibt. Der Nichtfachmann hat überhaupt keine andere Wahl, als dem Computer die Eigenschaften zuzuordnen, die durch die von der Presse verstärkten Propaganda der Computergemeinschaft zu ihm dringen. Daher hat der Informatiker die enorme Verantwortung, in seinen Ansprüchen bescheiden zu sein.“ (Ders., Alpträum Computer. Ist das menschliche Gehirn nur eine Maschine aus Fleisch? In: ZEIT 27 (1972) Nr. 3 v. 21. Januar 1972, 43.)

23 Vgl. Butlin, Patrick u.a., Consciousness in Artificial Intelligence: Insights from the Science of Consciousness, zu finden unter: <https://doi.org/10.48550/arXiv.2308.08708> (abgerufen am 28. Feb. 2028).

24 Vgl. Monya, Hannah u.a., Das Manifest. Elf führende Neurowissenschaftler über Gegenwart und Zukunft der Hirnforschung, in: Gehirn & Geist 3 (2004) Nr. 6, 30-37.

Qualia-Problems.²⁵ Die Frage, ob KI eines Tages menschenähnliche Fähigkeiten besitzen wird, lässt sich somit nicht rein technologisch beantworten, sondern erfordert eine philosophische Positionierung innerhalb der Mind-Brain-Debatte.²⁶

Sensibilität für unbeabsichtigtes Deskilling entwickeln

Wie bei fast allen Werkzeugen und technischen Hilfsmitteln ist ein Risiko der Nutzung von KI das „Deskilling“, also der **potenzielle Verlust menschlicher Kompetenzen** durch ihre unsachgemäße und unreflektierte Verwendung mit individuellen und kollektiven Folgen²⁷, auf den bereits erste Studien hinweisen.²⁸ Problematisch ist das, weil (a) KI-Systeme ausfallen können, (b) Menschen auch künftig in der Lage sein sollen, KI-Systeme zu überwachen und damit die jeweiligen Aufgaben auch selbst beherrschen müssen, und (c) durch die Auslagerung von Aufgaben und sozialen Interaktionen an KI ein **zunehmender Kontrollverlust** eintreten könnte.²⁹ Diskutiert wird das Phänomen vor allem im Kontext von Prüfungen (potenzielle Täuschungsversuche durch KI-Einsatz), wissenschaftlichem Arbeiten (Auslagerung von Schreib- und Denkprozessen an die KI) und Forschung (geringere Forschungsqualität, weniger Diversität und sinkende wissenschaftliche Integrität).³⁰ Zu unterscheiden sind die hier gemeinten unbeabsichtigten und letztlich unerwünschten Kompetenzverluste

25 Unter dem Qualia-Problem versteht man die Fragestellung nach der Verhältnisbestimmung von subjektiv-phänomenaler Wahrnehmung und mentalen Zuständen, vgl. auch Chalmers, David J., Facing Up to the Problem of Consciousness, in: Journal of Consciousness Studies 2 (1995) Nr. 3, 200-219.

26 Sowohl der ehemalige Google-Ingenieur Blake Lemoine, der dem Sprachmodell LaMDA ein Bewusstsein unterstellte, als auch der Futurist Ray Kurzweil antworten auf die Frage, warum man einer Maschine in Bezug auf ein behauptetes Bewusstsein glauben sollte, lediglich mit Glaubenssätzen: „benefit of the doubt“ (Lemoine) und „leap of faith“ (Kurzweil), vgl. Wan, Martin, Kann KI bewusst sein?

27 Vgl. Deutscher Ethikrat, Mensch und Maschine – Herausforderungen durch Künstliche Intelligenz, Stellungnahme v. 20. Mrz. 2023, zu finden unter: <https://www.ethikrat.org/fileadmin/Publikationen/Stellungnahmen/deutsch/stellungnahme-mensch-und-maschine.pdf> (abgerufen am 28. Feb. 2025), 353ff.

28 Vgl. Lee, Hao-Ping u.a., The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and Confidence Effects From a Survey of Knowledge Workers, zu finden unter: https://www.microsoft.com/en-us/research/wp-content/uploads/2025/01/lee_2025_ai_critical_thinking_survey.pdf (abgerufen am 28. Feb. 2025) sowie Gerlich, Michael, AI Tools in Society: Impacts on Cognitive Offloading and the Future of Critical Thinking, in: Societies 15 (2025) Nr. 6, zu finden unter: <https://doi.org/10.3390/soc15010006> (abgerufen am 28. Feb. 2025).

29 Vgl. Reinmann, Gabi, Deskilling durch Künstliche Intelligenz? Potenzielle Kompetenzverluste als Herausforderung für die Hochschuldidaktik (= HFD-Diskussionspapier Nr. 25 v. Okt. 2023), zu finden unter: https://hochschulforumdigitalisierung.de/sites/default/files/dateien/HFD_DP_25_Deskilling.pdf (abgerufen am 28. Feb. 2025), 7.

30 Vgl. ebd., 8.

von langfristigen Verschiebungen von Kompetenzprofilen, für die sich gute Gründe anführen lassen.

Vereinzelte gelangten Vorstöße wie der der Wirtschaftsuniversität Prag, angesichts von generativer KI Bachelorarbeiten abzuschaffen, in die mediale Aufmerksamkeit.³¹ Findet hier bereits eine solche Verschiebung von Kompetenzprofilen statt? Akademisches Schreiben ist als individueller Prozess der wissenschaftlichen Befassung mit einem Thema weiterhin ein wichtiger Teil des Kompetenzerwerbs für Studierende und der Qualifizierung in der Wissenschaft. Schriftliche Hausarbeiten und Qualifikationsschriften sollen die **Befähigung zum eigenständigen wissenschaftlichen Arbeiten** – inwiefern eine akademische Fragestellung in einer schriftlichen Darstellung wissenschaftlichen Standards genügend und originell bearbeitet werden kann – nachweisen; noch vor den Ergebnissen stehen die Dokumentation und Begründung des Untersuchungsdesigns, der gewählten Methoden und des individuellen Erkenntnisprozesses. Insofern lässt sich der Prager Vorstoß als Beispiel für möglicherweise unbeabsichtigtes Deskillung betrachten, sofern keine alternativen Formen gefunden werden, die den zuvor beschriebenen Kompetenzerwerb ähnlich effektiv ermöglichen.

Regeln für den Einsatz von KI in Studium, Lehre und Prüfungen seitens der Hochschule können einen verbindlichen Rahmen für individuelle Handlungen von Lehrenden und Studierenden schaffen. Ein solcher muss allerdings ausreichend diskutiert und verständlich gemacht werden. Die **Selbstbestimmung aller Akteure an der Hochschule** im Umgang mit generativer KI ist zu stärken, um so auch potenziellen Kompetenzverlusten entgegenzuwirken. Neben einem allgemeinen Verständnis der Technologie ist damit auch eine Haltung verbunden, sich nicht von der Technik kontrollieren und steuern zu lassen. Daneben ist zu diskutieren, welche **KI-unabhängigen Basiskompetenzen** vermittelt werden sollen. Das lässt sich nicht fachübergreifend bestimmen und ist eine genuin wissenschaftsdidaktische Herausforderung. Ein **humanzentriertes Zusammenspiel zwischen Mensch und Maschine** stärkt den Grundgedanken, dass sich technische Eigenschaften und menschliche Fähigkeiten komplementär ergänzen können – gepaart mit einer Sensibilität dafür, was der Mensch dabei verlernen kann und darf. So kann Bewusstsein für unbeabsichtigtes Deskillung geschaffen werden, ohne auf die Vorzüge von KI zu verzichten.³²

Rechtliches Grundverständnis aufbauen

Seit dem 2. Februar 2025 gelten die ersten Vorgaben der EU-KI-Verordnung. Hochschulen müssen seitdem sicherstellen, dass alle Akteure „rechtlich und technisch einschätzen

31 Vgl. Bager, Jo, ChatGPT & Co.: Uni in Prag schafft Bachelorarbeiten ab, in: Heise Online, zu finden unter: <https://www.heise.de/news/ChatGPT-Co-Uni-schafft-Bachelorarbeiten-ab-9546851.html> (abgerufen am 7. Mrz. 2025).

32 Vgl. Reinmann, Deskillung, 12 f.

können, was es bedeutet, KI einzusetzen.“³³ Fragen, die Hochschulen beantworten können müssen, sind: „**Was ist ein KI-System im Sinne der KI-VO?** Für wen gilt das KI-Recht? Welche Besonderheiten gelten für Hochschulen als öffentliche Stellen bei der Verwendung von KI? Wann ist KI erlaubt, aber hochriskant? Woran lassen sich hochriskante Anwendungen erkennen? Unter welchen Bedingungen darf hochriskante KI konkret verwendet werden? Wann ist KI verboten?“³⁴

Die Hochschulrektorenkonferenz wird hierzu Unterstützungsangebote und Handreichungen bereitstellen.

Atmosphäre des Misstrauens vermeiden

Vielfach wird im Kontext von generativer KI einseitig das Thema „Prüfungsbetrug“ in den Vordergrund gerückt. Dabei ist Prüfungsbetrug kein neues Phänomen und wird auch durch neue Technologien nicht zwangsläufig häufiger. **Hochschulen mussten immer schon darauf vertrauen, dass Hausarbeiten eigenständig und ohne die Hilfe weiterer Personen oder Ghost-writer erarbeitet wurden. Eine pauschale Misstrauenskultur gegenüber Studierenden wird dem Großteil jener, die KI redlich und kreativ einsetzen, nicht gerecht.** Stattdessen sollte eine konstruktive Einbindung digitaler Werkzeuge gefördert werden, die das bestehende Vertrauen zwischen Lehrenden und Lernenden stärkt und die Lernqualität verbessert.

In Ergänzung dazu sollten **Prüfungsformate** regelmäßig überprüft und, wenn nötig, angepasst werden. Kompetenzorientierte Aufgaben, die eine tiefere Auseinandersetzung mit den Lerninhalten anregen, Reflexions- und Transferleistungen, sind durch generative KI allein nur unzureichend zu lösen. Eine Kombination aus Vertrauen, klarer Kommunikation der Prüfungsanforderungen und gezielter didaktischer Integration von KI hilft, die Lernqualität nachhaltig zu stärken und Betrugsrisiken zu minimieren.

Visionen entwickeln, um reflektierte Entscheidungen treffen zu können

Die Reflexion der eigenen AG-Arbeit zum Ende hin hat uns immer deutlicher werden lassen, wie wichtig es bei Fragen zum Umgang mit KI im Hochschulkontext ist, die Fähigkeit zu haben, sich – zumindest zeitweise und zu geeigneten Zeitpunkten – von dem zu lösen, was faktisch ist, und argumentativ begründet über alternative Möglichkeiten nachzudenken. Gemeint ist die Entwicklung von Zukunftsszenarien und -entwürfen für Hochschulbildung mit KI. Der hier gewählte Weg, sich alternativen Möglichkeiten über die Erarbeitung konstruktiver Szenarien zu nähern, hat sich für uns – auch in der Diskussion mit Expert:innen – als überaus fruchtbar und als Verfahren erwiesen, **Orientierung zu schaffen und Handlungsmächtigkeit und -fähigkeit** da (wieder)herzustellen, wo man sich in Anbetracht der Wucht und Geschwindigkeit in der KI-Entwicklung allzu häufig ohnmächtig fühlt. Bildungs-

33 Schwartmann, Rolf, zit. nach: AA.VV., Erste Vorgaben zum Einsatz von KI greifen, in: Forschung und Lehre (online), zu finden unter: <https://www.forschung-und-lehre.de/recht/erste-vorgaben-zum-einsatz-von-ki-greifen-6878> (abgerufen am 28. Feb. 2025).

34 Ebd.

szenarien bis hin zu Gesellschaftsentwürfen entwickeln zu können – und dabei geht es nicht um weltfremde Hirngespinnste, sondern um **konstruktive, direkt nutzbare oder das Denken schärfende Visionen** –, könnte eine bislang vernachlässigte Kompetenz für einen geeigneten Umgang mit KI sein, um das zu überwinden, was Mittelstraß bereits im Kontext der Digitalisierung insgesamt am Bildungssystem und in der Gesellschaft kritisiert hat: nämlich dass der Fortschritt³⁵ nicht mehr uns gehört, sondern wir dem Fortschritt. Diese Fähigkeit jedoch scheint uns in großen Teilen der Gesellschaft und insbesondere in der Wissenschaft abhandengekommen zu sein. Schon 1964 benannte Adorno das als die „**seltsame Schrumpfung des utopischen Bewusstseins**“. Die aktuelle Vielzahl vor allem reaktiver, vermutlich nur kurzfristig wirksamer, mitunter sogar sachlich fragwürdiger und dann desorientierender (statt orientierender), vor allem aber fantasieloser Leitlinien und Handlungsempfehlungen zum Einsatz von KI³⁶ könnte ein Anzeichen sein, das Adornos Diagnose auch heute stützt.

Die Szenarienarbeit, wie wir sie für dieses Papier vorrangig betrieben haben, stellt *einen* Weg dar, sich der Aufgabe der Orientierung in Zeiten von Krisen wissenschaftlich und damit methodisch zu nähern. Eine weitere Möglichkeit ist der Übergang von der Konstruktion faktischer Szenarien hin zu kontrafaktischen Gedankenexperimenten.³⁷ Auch dieser Weg scheint vielversprechend, wobei er aufgrund seiner Kontrafaktizität methodisch noch voraussetzungsvoller ist. Die Fähigkeit des Denkens und Argumentierens in Möglichkeiten wiederherzustellen, scheint uns eine der zentralen Herausforderungen im Bildungssystem zu sein, um sich selbst orientieren und Gesellschaft gestalten zu können, und nicht lediglich auf Gestaltungsentwürfe kommerzieller oder ideologischer Akteure zu reagieren in Anbetracht aller möglichen Krisen. Nur so lassen sich **differenzierte und kontextabhängige Konzepte zum Umgang mit KI** entwickeln – jenseits naiver Euphorie und dystopischen Fatalismus.

Ein Denken und Argumentieren in Möglichkeiten geht weg von pauschalisiert-ideologischen „Ob-oder-ob-nicht“-Fragen hin zu „Wie“-Fragen, die es vor dem Hintergrund von Zielen und Werten zu diskutieren gilt, wie wir es in diesem Papier versucht haben. Wir laden dazu ein, unsere Szenarien als Ausgangspunkt für **eigene Auseinandersetzung mit Möglichkeitsräumen von KI** in den jeweils relevanten Konstellationen zu nutzen oder, affirmativ oder kritisch angeregt, eigene Szenarien zu entwerfen und diese als Diskussionsanlass zu nutzen. Gleichmaßen sind stets wechselwirkende Gegenrichtungen, paradoxe Nebeneffekte, also

35 Vgl. Mittelstraß, Jürgen, Bildung in einer Wissensgesellschaft, in: heiEDUCATION Journal 3 (2019), 21-36, 29, zu finden unter: <https://doi.org/10.17885/heiup.heied.2019.3.23942> (abgerufen am 11. Mrz. 2025).

36 Vgl. auch Weßels, Doris, Schluss mit absurden KI-Regeln, in: ZEIT 80 (2025) Nr. 3 v. 16. Jan. 2025, 31, zu finden unter: <https://www.zeit.de/2025/03/kuenstliche-intelligenz-studium-hochschulen-regeln> (abgerufen am 11. Mrz. 2025).

37 Vgl. Gehring, Petra, Rechtspolitische Bemessung möglicher gesellschaftlicher Gefahren digitaler Technologien? Zwei Gedankenexperimente mit anschließender Erwägung, in: Schreiber, Gerhard; Ohly, Lukas (Hg.), KI-Text. Diskurse über KI-Textgeneratoren, Berlin 2024, 355-359, zu finden unter: <https://doi.org/10.1515/9783111351490-022> (abgerufen am 11. Mrz. 2025), sowie Reinmann, Gabi, Hüter, Kümmerer, Vor- und? Eine Universität der Avatare: Ein Gedankenexperiment, in: Impact Free Nr. 61 (2025), zu finden unter: https://gabi-reinmann.de/wp-content/uploads/2025/01/Impact_Free_61.pdf (abgerufen am 11. Mrz. 2025).

Risiken aller Art, zu berücksichtigen und zu verhandeln, wie wir es in diesem Papier ebenfalls auf unterschiedliche Weise umgesetzt haben. Um der Komplexität von KI als einem zirkulären technologisch-gesellschaftlichen Gegenstand gerecht zu werden, bedarf es schlussendlich differenzierter interdisziplinärer Forschung, die neben einer empirischen Beschreibung konkreter Zusammenhänge auch gestaltend Möglichkeiten und Zukünfte entwerfen kann.

5 Literatur

AA.VV., Erste Vorgaben zum Einsatz von KI greifen, in: Forschung und Lehre (online), zu finden unter: <https://www.forschung-und-lehre.de/recht/erste-vorgaben-zum-einsatz-von-ki-greifen-6878> [abgerufen am 28. Feb. 2025].

Bager, Jo, ChatGPT & Co.: Uni in Prag schafft Bachelorarbeiten ab, in: Heise Online, zu finden unter: <https://www.heise.de/news/ChatGPT-Co-Uni-schafft-Bachelorarbeiten-ab-9546851.html> [abgerufen am 7. Mrz. 2025].

Butlin, Patrick u. a., Consciousness in Artificial Intelligence: Insights from the Science of Consciousness, zu finden unter: <https://doi.org/10.48550/arXiv.2308.08708> [abgerufen am 28. Feb. 2028].

Chalmers, David J., Facing Up to the Problem of Consciousness, in: Journal of Consciousness Studies 2 (1995) Nr. 3, 200-219.

Deci, Edward L.; Ryan, Richard M., Self-determination theory: A macrotheory of human motivation, development, and health, in: Canadian Psychology 49 (2008) Nr. 3, 182-185

Deutsche Forschungsgemeinschaft, Leitlinien zur Sicherung guter wissenschaftlicher Praxis. Kodex, Bonn 2022 (korrigierte Version 1.1), abgerufen am 21. Jan. 2025 unter: https://zenodo.org/records/6472827/files/kodex_leitlinien_gwp_dfg.1.1.pdf.

Deutscher Ethikrat, Mensch und Maschine – Herausforderungen durch Künstliche Intelligenz, Stellungnahme v. 20. Mrz. 2023, zu finden unter: <https://www.ethikrat.org/fileadmin/Publikationen/Stellungnahmen/deutsch/stellungnahme-mensch-und-maschine.pdf> [abgerufen am 28. Feb. 2025].

Ehlers, Ulf-Daniel, Future Skills. Lernen der Zukunft – Hochschule der Zukunft, Wiesbaden 2020.

Gehring, Petra, Rechtspolitische Bemessung möglicher gesellschaftlicher Gefahren digitaler Technologien? Zwei Gedankenexperimente mit anschließender Erwägung, in: Schreiber, Gerhard; Ohly, Lukas (Hg.), KI:Text. Diskurse über KI-Textgeneratoren, Berlin 2024, 355-359, zu finden unter: <https://doi.org/10.1515/9783111351490-022> [abgerufen am 11. Mrz. 2025].

Gerlich, Michael, AI Tools in Society: Impacts on Cognitive Offloading and the Future of Critical Thinking, in: Societies 15 (2025) Nr. 6, zu finden unter: <https://doi.org/10.3390/soc15010006> [abgerufen am 28. Feb. 2025].

Hanfeld, Michael, Die KI von Elon Musk verplappert sich, in: FAZ (online), zu finden unter: <https://www.faz.net/aktuell/feuilleton/grok-3-laesst-musk-und-trump-bei-frage-nach-fake-news-weg-110319792.html> [abgerufen am 28. Feb. 2025].

Huber, Ludwig, Hochschuldidaktik als Theorie der Bildung und Ausbildung, in: Ders. (Hg.), Ausbildung und Sozialisation in der Hochschule (= Enzyklopädie Erziehungswissenschaften

10), Stuttgart 1983, 127-129.

Jones, Nicola, OpenAI's 'deep research' tool: is it useful for scientists? Zu finden unter: <https://www.nature.com/articles/d41586-025-00377-9> [abgerufen am 28. Feb. 2025].

Ladenthin, Volker, Allgemeine Pädagogik, Baden-Baden 2022.

Lee, Hao-Ping u.a., The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and Confidence Effects From a Survey of Knowledge Workers, zu finden unter: https://www.microsoft.com/en-us/research/wp-content/uploads/2025/01/lee_2025_ai_critical_thinking_survey.pdf [abgerufen am 28. Feb. 2025].

Lee Myers, Steven, DeepSeek's Answers Include Chinese Propaganda, Researchers Say, in: New York Times [online], zu finden unter: <https://www.nytimes.com/2025/01/31/technology/deepseek-chinese-propaganda.html> [abgerufen am 28. Feb. 2025].

Major, Eddie, RAGs, Reasoning and Deep Research: What's new in AI and what might it mean for teaching in 2025? Zu finden unter: <https://www.adelaide.edu.au/learning/news/list/2025/02/21/rags-reasoning-and-deep-research-whats-new-in-ai-and-what-might-it-mean-for> [abgerufen am 28. Feb. 2025].

Miller, Riel, Futures Literacy. A hybrid strategic scenario method, in: Futures 39 [2007] Nr. 4, 341-662.

Mittelstraß, Jürgen, Bildung in einer Wissensgesellschaft, in: heiEDUCATION Journal 3 [2019], 21-36, 29, zu finden unter: <https://doi.org/10.17885/heiup.heied.2019.3.23942> [abgerufen am 11. Mrz. 2025].

Monya, Hannah u.a., Das Manifest. Elf führende Neurowissenschaftler über Gegenwart und Zukunft der Hirnforschung, in: Gehirn & Geist 3 [2004] Nr. 6, 30-37.

Nathan, Lisa P. u.a., Value scenarios. A technique for envisioning systemic effects of new technologies, in: CHI '07 Extended Abstracts on Human Factors in Computing Systems, San Jose 2007, 2585-2590, zu finden unter: <https://dl.acm.org/doi/10.1145/1240866.1241046> [abgerufen am 17. Feb. 2025].

Niesel, Dennis; Jelonnek, Stefan; Wilder, Nicolaus (in Druck), Gedankenexperimente als Methode pädagogischen Denkens – oder: Über die Notwendigkeit des Möglichen, Pädagogische Rundschau 2025.

Reinmann, Gabi, Deskillung durch Künstliche Intelligenz? Potenzielle Kompetenzverluste als Herausforderung für die Hochschuldidaktik (= HFD-Diskussionspapier Nr. 25 v. Okt. 2023), zu finden unter: https://hochschulforumdigitalisierung.de/sites/default/files/dateien/HFD_DP_25_Deskillung.pdf [abgerufen am 28. Feb. 2025], 7.

Dies., Gedankenexperimente als bildungstheoretisches Instrument in der Forschung zu Künstlicher Intelligenz im Hochschulkontext, in: Impact Free 56 [2014]. zu finden unter

https://gabi-reinmann.de/wp-content/uploads/2024/09/Impact_Free_58.pdf (abgerufen am 21. Jan, 2025).

Dies., Hüter, Kümmerer, Vormund? Eine Universität der Avatare: Ein Gedankenexperiment, in: Impact Free 61 [2025], zu finden unter: https://gabi-reinmann.de/wp-content/uploads/2025/01/Impact_Free_61.pdf (abgerufen am 11. Mrz. 2025).

Dies.; Watanabe, Alice, KI in der universitären Lehre, in: Schreiber, Gerhard; Ohly, Lukas (Hg.), KI: Text. Diskurse über KI-Textgeneratoren, Berlin 2024, 29-46.

Rosson, Mary Beth; Carrol John M., Scenario-based design, in: Sears, Andrew; Jacko, Julie A. (Hg.), The Human-Computer Interaction Handbook, New York 22008, 1041-1060; De Jouvanel, Hugues, A brief methodological guide to scenario building, in: Technical Forecasting and Social Change 65 [2000] Nr. 1, 37-48.

Se, Ksenia, How Do Agents Plan and Reason?, zu finden unter: <https://huggingface.co/blog/Kseniase/reasonplan> (abgerufen am 11. Mrz. 2025).

Wan, Martin, Kann KI bewusst sein? Anmerkungen zur Diskussion um LaMDA (Juli 2022), zu finden unter: <https://digiethics.org/2022/07/22/kann-ki-bewusst-sein-anmerkungen-zur-diskussion-um-lamda/> (abgerufen am 27. Feb. 2025).

Ders., Warum ChatGPT nicht das Ende des akademischen Schreibens bedeutet, zu finden unter: <https://digiethics.org/2023/01/03/warum-chatgpt-nicht-das-ende-des-akademischen-schreibens-bedeutet/> (abgerufen am 28. Feb. 2025).

Ders., Warum ChatGPT (auch weiterhin) nicht das Ende des akademischen Schreibens bedeutet, zu finden unter: <https://digiethics.org/2024/01/23/warum-chatgpt-auch-weiterhin-nicht-das-ende-des-akademischen-schreibens-bedeutet/> (abgerufen am 28. Feb. 2025).

Weizenbaum, Joseph, Altraum Computer. Ist das menschliche Gehirn nur eine Maschine aus Fleisch? In: ZEIT 27 [1972] Nr. 3 v. 21. Januar 1972, 43.

Weßels, Doris, Schluss mit absurden KI-Regeln, in: ZEIT 80 [2025] Nr. 3 v. 16. Jan. 2025, 31, zu finden unter: <https://www.zeit.de/2025/03/kuenstliche-intelligenz-studium-hochschulen-regeln> (abgerufen am 11. Mrz. 2025).

Wierda, Gerben, When ChatGPT summarises, it actually does nothing of the kind, zu finden unter: <https://ea.rna.nl/2024/05/27/when-chatgpt-summarises-it-actually-does-nothing-of-the-kind/> (abgerufen am 28. Feb. 2025).

Wilder, Nicolaus; Weßels, Doris, KI und das Ende des Einheitslehrplans? Eine kritische Analyse, in: Weiterbildung. Zeitschrift für Grundlagen, Praxis und Trends 35 [2025] Nr. 6.

Wissenschaftsrat, Empfehlungen zum Verhältnis von Hochschulbildung und Arbeitsmarkt. Zweiter Teil der Empfehlung zur Qualifizierung von Fachkräften (Drs. 4925-15), Bielefeld 2015, abgerufen am 21. Jan. 2025 unter: <https://www.wissenschaftsrat.de/download/archiv/4925-15.pdf>.

Impressum

Arbeitspapiere des HFD spiegeln die Meinung der jeweiligen Autor:innen wider.

Das HFD macht sich die in diesem Papier getätigten Aussagen daher nicht zu eigen.



Dieses Werk ist unter einer Creative Commons Lizenz vom Typ Namensnennung – Weitergabe unter gleichen Bedingungen 4.0 International zugänglich. Um eine Kopie dieser Lizenz einzusehen, konsultieren Sie <http://creativecommons.org/licenses/by-sa/4.0/>. Von dieser Lizenz ausgenommen sind Organisationslogos sowie, falls gekennzeichnet, einzelne Bilder und Visualisierungen.

ISSN (Online) 2365-7081; 11. Jahrgang

Zitierhinweis

Arbeitsgruppe „Künstliche Intelligenz: Essenzielle Kompetenzen an Hochschulen“ (2025). Künstliche Intelligenz: Grundlagen für das Handeln in der Hochschullehre. Arbeitspapier Nr. 86. Berlin: Hochschulforum Digitalisierung.

Herausgeber

Geschäftsstelle Hochschulforum Digitalisierung beim Stifterverband für die Deutsche Wissenschaft e.V.

Hauptstadtbüro • Pariser Platz 6 • 10117 Berlin • T 030 322982-520

info@hochschulforumdigitalisierung.de

Redaktion

AG „Künstliche Intelligenz: Essenzielle Kompetenzen an Hochschulen“

Verlag

Edition Stifterverband – Verwaltungsgesellschaft für Wissenschaftspflege mbH

Baedekerstraße 1 • 45128 Essen • T 0201 8401-0 • mail@stifterverband.de

Layout

Satz: Julia Rosche

Vorlage: TAU GmbH • Köpenicker Straße 154a • 10997 Berlin

Das Hochschulforum Digitalisierung ist ein gemeinsames Projekt des Stifterverbandes, des CHE Centrums für Hochschulentwicklung und der Hochschulrektorenkonferenz. Förderer ist das Bundesministerium für Bildung und Forschung.

www.hochschulforumdigitalisierung.de

