

Wir haben da ein ungutes Gefühl!

Anmerkungen zu generativer KI im Bildungskontext

Frank Vohle & Dominikus Herzberg

Abstract *Der Beitrag möchte auf ein ungutes Gefühl aufmerksam machen: Spätestens wenn generative Künstliche Intelligenzen in menschenähnliche Roboter verpackt werden und unsere Leben bevölkern, flechten sich die Maschinen in unsere Handlungs- und Beziehungsgestaltung ein. Wir leiten aus dieser Perspektive drei Problemebenen von generativer Künstlicher Intelligenz für Bildungsverhältnisse ab, nämlich real-unmittelbare, latent-emergente und spekulativ-erosive Probleme. Jede Ebene zeigt mögliche Risiken auf und stellt existenzielle Fragen. Die Betrachtungen zeigen, dass es einer Bildungsfolgenforschung bedarf, wenn einem daran gelegen ist, den Menschen im Blick zu behalten, wenn es um die Gestaltung von Bildung mit Künstlicher Intelligenz geht.*

Schlagwörter *Generative KI; humanoide Roboter; Risiko; Bildungsfolgenforschung; Bildungsphilosophie*

1. Einleitung: Wir gehen einem Gefühl nach

Bildung und Künstliche Intelligenz (KI) – bei allen Potenzialen, die wir sehen, überwiegt ein ungutes Gefühl. Das klingt unscharf und nebulös. Man könnte diese Gefühlslage wohlwollend als vorwissenschaftliche Ahnung bezeichnen. Wir sind uns dessen bewusst und möchten dieses ungute Gefühl dennoch geltend und zum Ausgangspunkt dieses Artikels machen: (1) Es geht um das Thema *generative* KI (im Text mit »GKI« abgekürzt), das den meisten im Bildungskontext in Form von Anwendungen begegnet, bei denen große Sprachmodelle (Large Language Models, LLMs) eine entscheidende Rolle spielen. Besondere Aufmerksamkeit gehört dabei Dialogpartnern wie ChatGPT, Gemini, Claude oder Deepseek. Die Technologie ist (trotz Vorläufern) radikal neu, sie

ist schlagartig in viele Tätigkeitsbereiche des Lebens eingedrungen und hat daher noch unübersehbare Implikationen für Bildung und Gesellschaft. Gefühle im Sinne von *leisen Ahnungen* können bei diesen unübersehbaren Implikationen als Such- und Erkenntnisressource hilfreich sein. (2) Wir beide sind seit jeweils ca. 20 Jahren in der digitalen Bildung professionell aktiv – der eine an der Hochschule, der andere in einem forschungsnahen Bildungsunternehmen. Wenn wir also von einem unguten Gefühl bei der Betrachtung von Möglichkeiten und Grenzen von GKI sprechen, meldet sich darin unsere *implizite Expertise* an, die wir zum kritischen Diskurs einbringen wollen. (3) Und schließlich: Gabi Reinmann hat uns über viele Jahre zum *mutigen Denken* angestiftet und herausgefordert. Grund genug, in dieser Festschrift einem vagen Gefühl zu folgen, was das Gegenteil einer sicheren Partie ist.

Der Beitrag gliedert sich in drei weitere Abschnitte: Im zweiten Abschnitt zeigen wir am Beispiel des »Fußball-Roboters«, wie sich GKI spielerisch und ganz handfest in der analogen Welt vermittelt. Für uns ist das eine motivierende Analogie, um auf grundlegende und potenzielle Problemfelder aufmerksam zu machen. Im dritten Abschnitt wollen wir diese Problemfelder exemplarisch beschreiben, um mögliche Dimensionen und Ausprägungen zu explizieren. Die Ergebnisse helfen uns im vierten Abschnitt, die Idee einer »Bildungsfolgenforschung« anzudeuten, die wir nicht nur, aber auch als normativ-anthropologische Disziplin auslegen wollen.

2. Fußballroboter: Spielerische KI agiert in der Welt

Kennen Sie die Weltmeisterschaft der Fußballroboter? Die »RoboCups« finden seit 1997 jährlich statt. Der Spaß daran, die eigenen Kreaturen als Team auf den Platz und in den sportlichen Wettkampf zu schicken, treibt rund 2.000 Wissenschaftler*innen und Studierende zusammen. Die Veranstaltung ist gleichzeitig Mittelpunkt einer wissenschaftlichen Konferenz. Das Ganze ist für alle Beteiligten nicht nur ein großes Spektakel, es stellt zudem eine wichtige und praktische Erkenntnisquelle für die Forschenden aus dem Bereich von KI und Robotik dar.

Mit dem RoboCup-Beispiel verbinden wir drei Beobachtungen, die nicht auf das Fußballspiel beschränkt bleiben sollen:

- 1) Die Verschmelzung von KI-Software mit Roboter-Hardware, wobei letztere eine menschenähnliche Gestalt aufweist.

- 2) Die Einbettung dieser neuen »Akteure« in realweltliche und zugleich spielerische Kontexte, die lange Zeit nur den Menschen vorbehalten waren.
- 3) Das Potenzial, dass Menschen nicht nur Roboter spielen »lassen«, sondern selbst als aktive Mitspieler auf dem Feld Teil des Spiels und eines Spielsystems werden, eine Art hybrides Spiel von Menschen und humanoiden Maschinen.

Der spielerische Kontext lässt die Beobachtungen unproblematisch erscheinen. Es fällt jedoch nicht schwer, sich das Training in den unzähligen Fußballvereinen in Deutschland zusammen mit ausgereiften humanoiden Fußballrobotern vorzustellen. Die Maschinen passen ihr Spielverhalten an das Leistungsniveau der Kinder und Jugendlichen an, unterstützen Trainer*in perfekt beim Einüben von Spielsituationen und rufen motivierende und antreibende Ansagen über den Platz, teils mit personalisierter Ansprache.

Das Beispiel deutet an, welche problematischen Implikationen sich beim Einsatz von KI-Robotern in der Bildung ergeben können: (a) Durch die menschenähnliche Gestalt (und Funktion) wird es in Zukunft emotional schwieriger, die Grenze zwischen einem KI-Roboter und einem Menschen zu markieren, was Folgen für das Selbstverhältnis des Menschen und der Beziehungsqualität zwischen den Menschen hat. (b) Durch die »spielerische« Einführung von KI-Robotern wird jener Rückzugsort belegt, der insbesondere für das Mensch-Sein und Mensch-Werden attraktiv ist, also neben Spiel auch Kunst, Poesie etc. (c) Durch das Potenzial der Hybridisierung (*humanoid-roboter-hybrid-games*) wird die für Menschen flüchtige, aber einmalige Situation (des Spiels) mit technischen Mitspielern geteilt, die darin lediglich eine Lerngelegenheit identifizieren, die sie mit nichtanwesenden Ko-Akteuren in einem globalen Informationsnetz in Echtzeit zum Zweck der Leistungsoptimierung »verrechnen«; hier wird gemeinsame Sinnkonstruktion lediglich simuliert.

Dieses Einstiegsbeispiel ist weniger spekulativ, als man meinen mag. Das Verhältnis von Mensch und Maschine »vermenschlicht« sich über natürlich-sprachliche Interaktionen. So kombiniert etwa die Firma Figure die ChatGPT-Sprachmodelle mit humanoiden Robotern. Mit dem Roboter »Figure 02« liegt ein Prototyp vor, der zeigt: Wir reden bei GKI nicht mehr nur von intelligenten Dialogpartnern, mit denen wir via Laptop oder Handy »mental in Kontakt kommen«, sondern von »intelligenten« Akteuren innerhalb unserer physikalisch-analogen Welt. Das hat Folgen für die Interpretation von GKI im Allgemeinen und für das Bildungssystem im Besonderen.

3. Ausgewählte Problemfelder bzw. -levels

Das Einstiegsbeispiel hat vor allem den Zweck, Sie als Leser*in aufmerksam zu machen für *mögliche* und vor allem auch *langfristige* Entwicklungen. Damit reden wir nicht gegen die unübersehbaren Chancen von GKI und Robotik, zum Beispiel für Pflege, Sicherheit, Medizin und Wissenschaft. Aber: Wir möchten im Folgenden mit unserem unguuten Gefühl den Blick auf die »dunkle Seite« richten, um einen Beitrag zur Exploration des Feldes zu leisten, was für jedes aufgeklärte Handeln unerlässlich ist bzw. sein sollte!

Hierzu schlagen wir drei Problemfelder vor, die wir durch Beispiele verdeutlichen und in ihren Merkmalen charakterisieren. Das für uns Interessante ist, dass die Problemfelder von bekannten, im KI-Alltag anzutreffenden Problemen ausgehen und zunehmend spekulativer, aber aus Bildungssicht kritischer, riskanter und existentieller werden, weswegen wir auch von »Problem-levels« sprechen. Wir orientieren uns dabei vorrangig an GKI, da sie derzeit die eindringlichste KI-Technologie ist, die jeder Person im Büro- und Alltagsbereich wie auch im Bildungsbereich zur Verfügung steht und gesellschaftliche Debatten ausgelöst hat. Unsere Beispiele lösen sich vom Zugang über humanoide KI-Roboter und adressieren sehr allgemein den Bildungsbereich im Kontext von GKI.

3.1. Level 1: Real-unmittelbare Probleme

Auf dem ersten Level liegen Probleme, die *empirisch gut dokumentiert* sind und seit einiger Zeit diskutiert werden, zum Beispiel Halluzination und Deep Fakes.

Als »Halluzination« bezeichnet man das Problem KI-gestützter Sprachmodelle, dass Ausgabeergebnisse nicht wahr sein müssen, sondern vom System frei erfunden und damit unsinnig sein können. Diese Eigenschaft ist in der Natur generativer Sprachmodelle angelegt, die lediglich wahrscheinliche Sprachfortsetzungen hervorbringen, die sich aus dem Training mit Texten ergeben haben. Das funktioniert zwar erstaunlich gut und erzeugt oft anschlussfähige Kommunikation, aber eine Garantie für die Faktizität einer Ausgabe kann nicht gegeben werden.

Als »Deep Fakes« bezeichnet man realistisch wirkende Medieninhalte, zum Beispiel Fotos oder Videos, die durch GKI erzeugt oder verändert worden sind. Über soziale Medien lassen sich Falschnachrichten mit realistisch wirkendem Bild- oder Videomaterial verbreiten, die Menschen in Szenarien einbetten und

Dinge sagen und tun lassen, die erfunden sind. Die Glaubwürdigkeit von Medieninhalten erodiert, das Missbrauchspotenzial ist groß – und wird genutzt.

Die Beispiele mögen genügen, um das Problemlevel 1 zu charakterisieren: Es sind Probleme, die (a) *bekannt* und empirisch nachweisbar sind, die (b) hinsichtlich ihres Entstehungs- und Wirkungszusammenhangs weitestgehend *verstanden* sind, die man (c) aus normativen Überlegungen für *nicht wünschenswert* einstuft und (d) an deren Lösung man aufgrund der *Akzeptanz von a-c* mit unterschiedlichen Mitteln arbeitet.

3.2. Level 2: Latent-emergente Probleme

Auf dem zweiten Level liegen Probleme, die aktuell nur randständig diskutiert werden, die über anekdotische Evidenz hinaus schwer zu fassen, jedoch gut *logisch zu verargumentieren* sind. Zu diesen Problemen zählen das »Deskilling« oder die »Asozialität«.

Unter »Deskilling« versteht man den Kompetenzverlust mit individuellen und kollektiven Folgen, der gerade durch den Einsatz von KI im Bildungsbereich möglich ist (Reinmann, 2023b). Vor allem generative KI nimmt Einfluss auf das, was viele Autor*innen als »Wissensarbeit« bezeichnen. Damit wird angedeutet: Es geht bei der generativen KI nicht lediglich um ein neues Werkzeug, das zum Beispiel bei körperlich oder geistig anspruchsvollen Prozessen entlasten soll, sondern es geht potenziell um eine *unspezifische und unkontrollierte Entlastung des Denk- und Sprachkerns des Menschen, ohne Aussicht auf höherwertigen Ersatz*. Gerade darin wird aber die Bedingung der Möglichkeit zu einem grundlegenden Kompetenzverlust gesehen: Unspezifisch ist die Entlastung, wenn zum Beispiel Lehrende kein Feedback mehr geben, weil das ein KI-gestütztes Feedbacksystem macht. Unkontrolliert ist die Entlastung, wenn aus Zeit- und Ressourcenmangel oder schlichter Bequemlichkeit »fast alles« durch GKI organisiert, vermittelt, angeeignet, begleitet und geprüft wird und damit auch die Zielsetzungs- und Kontrollautonomie des Menschen abgegeben wird – freiwillig!

Von »Asozialität« spricht Herzberg (2024) dann, wenn in einer GKI-gestützten Kommunikation (z. B. E-Mail, Hausarbeit) die soziale Beziehung dann gestört wird, wenn bisher kongruente Erwartungen zum Beispiel zur Eigenleistung bei der Texterstellung einseitig und intransparent gekündigt und damit in inkongruente Erwartungen überführt werden. Im schlimmsten Fall erodiert die soziale Beziehung hin zur Asozialität. In ähnlicher Weise spricht Kompa (2024) davon, dass menschliche Emotionen in den KI-ge-

stützten Medien als »mediale Mikrosensationen inszeniert« seien, dass die jüngere Generation der Digital Natives »in der Illusion aufwache, mit Quasi-Menschen zu sprechen« mit der Konsequenz, dass dort, wo keine soziale Resonanz benötigt wird, es keine Entwicklung von Beziehungsfähigkeit (Welt-, Freundschafts-, Gemeinschafts- und Liebesbeziehungen) gebe (S. 3).

Problemlevel 2 zeichnet sich durch Probleme aus, die (a) *wenig bekannt* sind, weil sie mit einer Verzögerung in der Zukunft eintreten können, die (b) hinsichtlich ihres Entstehungs- und Wirkungszusammenhangs weitestgehend *unverstanden* sind, die man (c) aus normativen Überlegungen für *nicht wünschenswert einstuft* und (d) deren Lösung man aufgrund des offenen Diskussionsstandes von a-c aktuell wenig Beachtung schenkt.

3.3. Level 3: Spekulativ-erosive Probleme

Auf dem dritten Level liegen Probleme, die im Bildungskontext noch nur sehr selten diskutiert werden, aber die individuellen und kollektiven *Voraussetzungen von Bildung* zerstören. Hierzu zählen Formen von Freiheitsverlust und totalitäre Zwänge.

Als »freiwilligen Freiheitsverlust« sollen alle Aktivitäten bezeichnet werden, die Menschen – vor allem jene, die nicht ausreichend lesen, schreiben und rechnen können – *ohne Zwang*, aber durch subtile psychologische Tricks (z. B. Nudging) meist lebenslang an KI binden: angefangen von Reiseplanung und Restaurantauswahl über E-Learning-Angebote und Bewerbungsschreiben bis hin zu Partner-, Politik- und Finanzentscheidungen. Kompa (2024) spricht in diesem Zusammenhang von einem »*relativen* Freiheitsverlust, ganz im Sinne Kants, durch Abtretung einer offenen Selbstbestimmung an externe Systeme« (S. 2)

Wir sprechen von »totalitär ähnlichen Zwängen« (Watanabe, 2024), wenn die Handlungsalternativen ganzer Gesellschaften nach der *Struktur des sozialen Dilemmas* (Beckenkamp & Vohle, 2022) organisiert sind, das heißt wenn es zwar individuell sinnvoll ist, digitale Angebote zu nutzen (günstig, mehrwertig, sicher, schnell etc.), aber genau in dieser Nutzung ein kollektiv-kultureller Schaden entsteht, der das Fundament freiheitlich demokratischer Gesellschaften angreift. Das können zum *einen sprach- und denknormierende Effekte* von wenigen, aber global bestimmenden KI-Systemen sein oder die Entwicklung (*quasi-*)*monopolistischer* KI-Produkte, deren *Nichtnutzung zwar möglich, aber mit großen individuellen Nachteilen verbunden* ist. Zum Beispiel kann die Verweigerung von KI-gestützten Bildungs-, Mobilitäts- oder Gesundheitsprogrammen

dazu führen, dass man nicht von geeigneten Angeboten oder passenden Beiträgen profitieren kann (Peetz, 2024).

Das Problemlevel 3 kennzeichnet Folgendes: Es sind Probleme, die (a) *wenig bekannt* sind, weil sie mit einer Verzögerung in der Zukunft eintreten können, die (b) hinsichtlich ihres Entstehungs- und Wirkungszusammenhangs weitestgehend noch *unverstanden* sind und zudem mit *ideologischen und damit wissenschaftlich fragwürdigen Annahmen* arbeiten, die man (c) aus normativen Überlegungen für *nicht wünschenswert einstuft* und (d) deren Lösung man aufgrund des offenen Diskussionsstandes *von a-c* aktuell wenig Beachtung schenkt.

Wir haben nun drei Problemlevels mit je zwei Beispielen knapp beschrieben. Die Problemlevels sind mit den Begriffen »real-unmittelbar«, »latent-emergent« und »spekulativ-erosiv« gekennzeichnet und die Beispiele so ausgewählt, dass darin der Charakter der Levels anschaulich wird. Insgesamt hatte das Kapitel den Sinn, das Problemfeld in einem *ersten Anlauf zu umreißen*, um mögliche Perspektiven und Kategorien zu unserem »ungenuten Gefühl« zu explizieren.

4. Bildungsfolgenforschung

Seit den 1960er Jahren gibt es in der Technikphilosophie das Forschungsgebiet der Technikfolgenabschätzung. Ziel dieser Disziplin ist es, sowohl Fakten-, Ursachen- und Prognosewissen als auch Orientierungs- und Handlungswissen zu Wissenschaft und Technik für verschiedene Akteure, zum Beispiel in der Politik, bereitzustellen (Saretzki, 2014). Für die Bildung gibt es ein solches Konzept bislang nicht!

Angesichts der fortschreitenden Technisierung der Bildung, insbesondere durch GKI, fordert Herzberg (2023) eine »Bildungsfolgenforschung«. Eine solche Idee erhält Rückenwind durch die neuen Herausforderungen, die in den zuvor skizzierten Problemfeldern sichtbar wurden. Wenn wir die »digitale Souveränität« von Individuen und Kollektiven in Bildungskontexten zumindest in der EU-Wertegemeinschaft aufbauen und verteidigen möchten (Council of Europe, 2024), dann wäre die Bildungsfolgenforschung im Hochschulkontext ein Lösungsbaustein. Wir möchten abschließend ein paar Impulse zur Diskussion stellen, wie wir uns eine Bildungsfolgenforschung vorstellen und welche Merkmale sie haben sollte:

- 1) Multiparadigmatik! Zum ersten könnte sich eine Bildungsfolgenforschung von den Perspektiven inspirieren lassen, die im Abschnitt drei skizziert wurden. Die Problemlevels »real-unmittelbar«, »latent-emergent« und »spekulativ-erosiv« haben sehr verschiedene Reichweiten, Eintrittswahrscheinlichkeiten und Implikationen. Sie lassen sich am besten im transdisziplinären Verbund bearbeiten und mit unterschiedlichen paradigmatischen Zugängen lösen. So sollten neben den aktuell gängigen empirischen Zugängen zum Beispiel auch theoretisch-argumentative Verfahren (z. B. das philosophische Gedankenexperiment, Reinmann, 2024) oder auch entwicklungs- und designorientierte Beiträge (s. u.) genutzt werden. Die Bildungsfolgenforschung wird hier auf ähnliche Kritik stoßen wie die Technikfolgenabschätzung, nämlich die, dass das Finden von geeigneten Expert*innen schwierig und die Generierung von Zukunftswissen generell angreifbar ist.
- 2) Transdisziplinarität! Zum zweiten sollte die Bildungsfolgenforschung keine reine akademische Aktivität im Elfenbeinturm sein. Mit dem Begriff der »Transdisziplinarität« sollte angedeutet werden, dass neben Wissenschaftler*innen aus den Hochschulen auch Vertreter*innen der Bildungspraxis aus Hochschulen, Schulen und Arbeitswelt sowie interessierte Bürger*innen zusammenarbeiten sollen (Vohle, Sygusch & Söhngen, 2024). In dieser »Open Science Community« kommen Expert*innen und Anwender*innen aus sehr unterschiedlichen Lebensbereichen zu Wort (Partizipation) und es vernetzen sich unterschiedliche Hochschulstandorte zum Beispiel zu einem »Forschungslabor« für bildungsnahe GKI.
- 3) Entwurforschung! Zum dritten möchten wir für eine Bildungsfolgenforschung explizit die Entwurfs- oder Designforschung in Stellung bringen (Reinmann, Herzberg & Brase, 2024), weil sie mögliche Bildungswelten entwickeln hilft und dabei spezifische Formen von Erkenntnis (unter anderem intentionale Kausalität, Bakker et al., 2024) freisetzt. Eine solche Forschung erlaubt es zwar nicht, die Komplexität einer Bildungszukunft mit GKI vorherzusagen, sie versetzt aber Teilnehmende der Forschung in die Lage, die komplexe Gemengelage aus ethischen Vorentscheidungen, soziotechnische Prototypen und kontextuellen Bedingungen zu situieren und einen (kritischen) Reflexionsrahmen zu erarbeiten. Das hilft Lernen- den, Lehrenden, Forschenden, Hochschulleitungen und schließlich auch Politiker*innen dabei, einsichtsgeleitete Fragen zu stellen.

Wir haben in der Einleitung davon gesprochen, dass sich eine Bildungsfolgenforschung auch als normativ-anthropologische Disziplin versteht bzw. verstehen muss. Wenn Reinmann (2023a) für den Bereich der Hochschule danach fragt: »Wozu sind wir hier?«, dann deutet sie damit die radikale Einbindung auch ethischer, normativer, letztlich auch anthropologischer Wissenschaften an. Während ethische Positionen und Handlungsempfehlungen gerade in der Technikfolgenabschätzung schon differenziert beschrieben sind (Lenk, 2009), ist die Frage der GKI in der Bildung nicht nur eine Frage der Verantwortungsethik, sondern auch eine der Anthropologie. Angesichts der potenziellen »Kopplung« von Menschen und GKI oder der »Koexistenz« humanoider Roboter mit Menschen in Arbeit, Bildung und Spiel steht zunehmend die Frage im Raum, wie »Menschsein« definiert ist: Wo liegen die Grenzen, zum Beispiel einer Entlastung durch GKI? Wie ändert sich das Selbstverständnis und Selbstverhältnis des Menschen, wenn neben ihnen, den Menschen, intelligente und ethisch versierte »Akteure« auftreten, die zwar kein biologisches Bewusstsein haben, aber intentional zu handeln *scheinen*, was sie in die Nähe von »Würdeträgern« bringt (Gloy, 2024, Knell, 2022). Diese Grenzarbeit lässt sich nicht mit empirischen Mitteln allein bewerkstelligen, sondern bedarf auch eines normativ-philosophischen Diskurses und regulierender (politischer) Institutionen wie dem des »Council of Europe«, zu dessen Weiterentwicklung die Bildungsfolgenforschung einen Beitrag leisten könnte.

Literatur

- Bakker, Arthur, Angerer, Elisabeth, Penuel, William & Akkerman, Sanne (2024). Causal Reasoning about education: What is it and what should it be? I Phyllis Illari & Federica Russo (Hg.), *The Routledge Handbook of Causality and Causal Methods* (pp. 671–682).
- Beckenkamp, M. & Vohle, P. (2022). Freiwillige Institutionen zur Überwindung der Tragödie der Allmende im Lichte erweiterter begrenzter Rationalität?! *Zeitschrift für Gemeinwirtschaft und Gemeinwohl*, 45(2), 377 – 396.
- Council of Europe (2024). *Artificial Intelligence and Education*. <https://rm.coe.int/dgii-edu-aied-2024-03-2nd-working-conference-concept-note-and-draft-pr/1680b0cd23> (21.10.2025).
- Gloy, Karen (2024). *Mensch oder Roboter, Natur oder KI?* Königshausen & Neumann.

- Herzberg, Dominikus (2023). Künstliche Intelligenz in der Hochschulbildung und das Transparenzproblem. In Tobias Schmohl, Alice Watanabe & Kathrin Schelling (Hg.), *Künstliche Intelligenz in der Hochschulbildung. Chancen und Grenzen des KI-gestützten Lernens und Lehrens*. transcript.
- Herzberg, Dominikus (2024). *Generative KI, Vollzüge in Sprache, Gedanken zur Bildung*. <https://www.youtube.com/watch?v=MKGp2kkFlzs> (18.12.2025).
- Knell, Sebastian (2022). Künstliche Intelligenz und menschliche Würde – ein aporetisches Verhältnis? *ZEMO*, 5, 203–229.
- Kompa, Joana (2024). *Die Kolonialisierung der Zukunft – eine Kritik*. <https://joana.kompa.com/2024/05/19/die-kolonialisierung-der-zukunft/> (21.10.2025).
- Lenk, Hans (2009). *Das flexible Vielfachwesen*. Velbrück.
- Peez, Thorsten (2024). Kapitalismus, digital. [Rezension zu: Tanja Carstensen, Simon Schaupp, Sebastian Seignani (Hg.), *Theorien des digitalen Kapitalismus*. Suhrkamp].
- Reinmann, Gabi (2023a). Wozu sind wir hier? Eine wertebasierte Reflexion und Diskussion zu ChatGPT in der Hochschullehre. *Impact Free*, 51, 1–12.
- Reinmann, Gabi (2023b). *Deskilling durch Künstliche Intelligenz?* Diskussionspapier Nr. 25. Hochschulforum Digitalisierung.
- Reinmann, Gabi (2024). Gedankenexperimente als bildungstheoretisches Instrument in der Forschung zu Künstlicher Intelligenz im Hochschulkontext. *Impact Free*, 58, 1–11.
- Reinmann, Gabi, Herzberg, Dominikus & Brase, Alexa (2024). *Forschendes Entwerfen*. transcript.
- Saretzki, Thomas (2014). Entstehung und Status der Technikfolgenabschätzung. *APuZ* 6–7/2014, 11–16.
- Vohle, Frank, Sygusch, Ralf & Söhngen, Markus (2024). Wissenschaftsdidaktik »abroad«. Eine Fallstudie zur transdisziplinären Kommunikation im Kontext Sport. In Gabi Reinmann & Rüdiger Rhein (Hg.), *Wissenschaftsdidaktik IV*. transcript.
- Watanabe, Alice (2024). *Forschen in Fragmenten. Dissertationsschrift*. Universität Hamburg.

