

KI:Medien:Ethik

Hochschullehrende zwischen Chancen, Befürchtungen und Verantwortung

*Sonja Gabriel¹, Elisabeth Zissler², Jennifer Jakob³,
Michaela Liebhart-Gundacker⁴*

DOI: <https://doi.org/10.53349/resource.2026.i3.a1579>

Zusammenfassung

Der Beitrag stellt Ergebnisse des quantitativen Forschungsteils des Projekts KI:Medien:Ethik vor. Im Mittelpunkt stehen Einstellungen von Hochschullehrenden an einer Pädagogischen Hochschule zu generativer Künstlicher Intelligenz sowie zu ausgewählten ethischen Fragestellungen. Grundlage bilden zwei Online-Erhebungen aus den Jahren 2024 und 2025, die deskriptiv vergleichend ausgewertet werden. Analysiert wird, welche Chancen, Befürchtungen und ethischen Problemwahrnehmungen Hochschullehrende mit KI-Systemen verbinden und wie sich Antworttendenzen zwischen den beiden Erhebungszeitpunkten verändern. Die Auswertung wird entlang ethischer Prinzipien wie Gerechtigkeit, Nicht-Nachteiligkeit, Vorteilhaftigkeit, Nachvollziehbarkeit, Autonomie und Privatsphäre strukturiert. Die Ergebnisse zeigen eine ausgeprägte Sensibilität gegenüber Risiken generativer KI, insbesondere in Bezug auf Datenschutz, Stereotypisierung, demokratische Prozesse und Nachvollziehbarkeit. Kausale Aussagen über Wirkungen hochschulischer Maßnahmen können auf Grundlage des vorliegenden Designs nicht getroffen werden. Die Ergebnisse zeigen, dass unter den befragten Hochschullehrenden eine hohe Sensibilität für ethische Risiken generativer KI besteht. Zugleich verweisen sie auf den Bedarf an institutioneller Orientierung, Fortbildung und diskursiven Formaten.

Stichwörter: Ethische Zugänge zu KI; Einstellungen, Ängste und Chancen gegenüber KI; Hochschulische AI-Literacy

¹ Kirchliche Pädagogische Hochschule Wien/NÖ, Mayerweckstraße 1, 1210 Wien.
E-Mail: sonja.gabriel@kphvie.ac.at

² Kirchliche Pädagogische Hochschule Wien/NÖ, Mayerweckstraße 1, 1210 Wien.
E-Mail: elisabeth.zissler@kphvie.ac.at

³ Kirchliche Pädagogische Hochschule Wien/NÖ, Mayerweckstraße 1, 1210 Wien.
E-Mail: jennifer.jakob@kphvie.ac.at

⁴ Kirchliche Pädagogische Hochschule Wien/NÖ, Mayerweckstraße 1, 1210 Wien.
E-Mail: michaela.liebhart@kphvie.ac.at

1 Einleitung

Die rasche Entwicklung generativer Künstlicher Intelligenz in den letzten Jahren stellt Hochschulen nicht nur vor technologische, sondern vor grundlegende normative Herausforderungen (Nguyen, 2025; Reinmann et al., 2025; Zhai, 2024). Während KI-Anwendungen zunehmend in Lehr-, Lern- und Forschungsprozesse integriert werden (Bosse et al., 2026; Prilop et al., 2024; Zhai, 2024), verschieben sich zugleich die Maßstäbe dafür, was als Wissen, Leistung und Verantwortung gilt. Hochschulen sind damit nicht nur Orte der Innovation, sondern auch zentrale Räume der Reflexion darüber, wie technologische Entwicklungen gesellschaftlich eingeordnet und gestaltet werden können.

Vor diesem Hintergrund kommt Hochschullehrenden eine Schlüsselrolle zu. Sie agieren nicht nur als Anwender*innen von KI, sondern als Vermittler*innen zwischen Technologie, didaktischer Praxis und gesellschaftlicher Verantwortung. Ihre Einstellungen, Wahrnehmungen und Unsicherheiten im Umgang mit (generativer) KI sind daher von zentraler Bedeutung für die Frage, wie sich Hochschulbildung im Kontext dieser Technologien weiterentwickelt. Insbesondere die ethische Dimension – etwa im Hinblick auf Gerechtigkeit, Transparenz, Autonomie oder Datenschutz – gewinnt dabei zunehmend an Gewicht (Floridi et al., 2018; Gaisch & Mader, 2025; Reinmann et al., 2025), da sie die Grundlage für eine verantwortungsvolle Integration von KI in Bildungsprozesse bildet.

Das Forschungsprojekt KI:Medien:Ethik (KI:M:E) setzt genau an dieser Schnittstelle an. Es untersucht, wie Hochschullehrende generative KI wahrnehmen, welche Chancen und Befürchtungen sie damit verbinden und inwiefern ethische Fragestellungen ihre Perspektiven prägen. Ergänzend wird deskriptiv erhoben, welche hochschulischen Angebote zum Thema KI von den Befragten wahrgenommen oder besucht wurden. Eine kausale Überprüfung der Wirkung dieser Angebote ist mit dem vorliegenden Design nicht verbunden.

Der vorliegende Beitrag fokussiert auf die Ergebnisse der quantitativen Erhebung innerhalb dieses Projekts. Ziel ist es, Einstellungen von Hochschullehrenden zu zentralen ethischen Fragestellungen systematisch zu erfassen, Antworttendenzen der beiden Erhebungszeitpunkte deskriptiv zu vergleichen und diese vor dem Hintergrund etablierter ethischer Prinzipien einzuordnen. Damit leistet der Beitrag einen empirisch fundierten Beitrag zur Frage, wie sich ethische Reflexionsprozesse im Kontext generativer KI in der Hochschullehre gestalten und weiterentwickeln lassen.

2 Hochschulen und KI

Die aktuelle Forschung zur Integration von (generativer) künstlicher Intelligenz in Hochschulen zeigt deutlich, dass ethische Fragen mittlerweile ebenfalls im Fokus der Auseinandersetzung stehen. KI wird nicht nur als didaktisches Werkzeug betrachtet, sondern als soziotechnisches System, das tief in Fragen von Gerechtigkeit, Verantwortung und Machtverhältnissen in

Bildung eingreift. Der Bereich der KI-Ethik beschäftigt sich grundsätzlich mit Prinzipien und Herausforderungen sowohl der KI-Entwicklung als auch deren Anwendung (Gaisch & Mader, 2025). Allerdings gibt es, bezogen auf den Hochschulbereich, auch hier noch größere Forschungsdefizite (Reinmann et al., 2025). So werden ethische Bedenken häufig als Begründung angeführt, warum (generative) KI in der Lehre nicht eingesetzt wird (Bosse et al., 2026).

Ein besonders prominenter Themenbereich, der mittlerweile gut erforscht ist, ist algorithmische Verzerrung und Diskriminierung. Studien zeigen, dass Hochschullehrende die Gefahr sehen, dass KI bestehende gesellschaftliche Ungleichheiten reproduziert oder sogar verstärkt. Dies betrifft sowohl Inhalte als auch Bewertungen und Interaktionen, da Trainingsdaten häufig kulturelle und soziale Bias enthalten (Ferreira et al., 2025; Medeiros et al., 2025). In diesem Zusammenhang wird auch die Dominanz westlich geprägter Daten und Perspektiven kritisch diskutiert, die zu einer eingeschränkten Repräsentation globaler Diversität führen kann.

Eng damit verbunden ist die Frage der Bildungsgerechtigkeit und des Zugangs. Der Einsatz von KI kann bestehende digitale Ungleichheiten verschärfen, wenn unterschiedliche Gruppen ungleichen Zugang zu Technologien, Infrastruktur oder Kompetenzen haben. Forschungsergebnisse weisen darauf hin, dass insbesondere sozioökonomische Unterschiede eine zentrale Rolle spielen und sich in einer neuen Form digitaler Ungleichheit manifestieren können (Suchismita, 2025; Tunjera et al., 2025).

Ein weiterer zentraler Aspekt ist die Transparenz und Nachvollziehbarkeit von KI-Systemen. Viele Anwendungen funktionieren als Black Box, wodurch Entscheidungsprozesse für Lehrende und Studierende schwer verständlich bleiben. Diese Intransparenz steht im Widerspruch zu zentralen Prinzipien wissenschaftlicher Bildung, die auf Nachvollziehbarkeit, Begründbarkeit und kritischer Reflexion beruhen (Macgilchrist et al., 2025; Medeiros et al., 2025).

Auch Datenschutz und Datensicherheit werden als zentrale ethische Herausforderungen beschrieben. Hochschulen arbeiten mit sensiblen personenbezogenen Daten, deren Verarbeitung durch KI-Systeme neue Risiken mit sich bringt. Lehrende äußern insbesondere Unsicherheit darüber, welche Daten erhoben werden, wie diese genutzt werden und wer Zugriff darauf hat (Ferreira et al., 2025).

Darüber hinaus steht die akademische Integrität im Fokus der Forschung (Nguyen, 2025). Generative KI stellt traditionelle Leistungsnachweise infrage, da Studierende Aufgaben mithilfe von KI lösen können. Dies führt zu grundlegenden Diskussionen über Authentizität, Eigenleistung und die Validität von Prüfungsformaten (Estaiteyeh et al., 2024; Petraki, 2024). Gleichzeitig wird betont, dass diese Herausforderung auch eine Chance darstellt, Prüfungsformate stärker auf Reflexion, Anwendung und kritisches Denken auszurichten.

Ein übergreifendes Problem ist das Fehlen klarer ethischer Leitlinien und institutioneller Rahmenbedingungen. Viele Hochschullehrende in unterschiedlichen Befragungen berichten von Unsicherheiten im Umgang mit KI, da verbindliche Richtlinien fehlen oder noch in Entwicklung sind. Dies betrifft Fragen wie den angemessenen Einsatz von KI, den Umgang mit studentischen KI-Nutzungen sowie die Verantwortung von Lehrenden (Ceallaigh et al., 2025; Ferreira et al., 2025).

Neben diesen Herausforderungen betont die Forschung die Notwendigkeit, ethische Reflexion systematisch in die Hochschullehre zu integrieren. Dies umfasst die Entwicklung von AI Literacy, die nicht nur technische Kompetenzen, sondern auch kritische, ethische und gesellschaftliche Perspektiven einschließt (Meylani, 2024; Tunjera et al., 2025). Lehrende übernehmen dabei eine Schlüsselrolle, indem sie Studierende befähigen, KI nicht nur zu nutzen, sondern deren Auswirkungen auf Individuum und Gesellschaft zu verstehen und zu hinterfragen.

Andere Aspekte, wie didaktische Potenziale, Effizienzgewinne oder Veränderungen von Lehrrollen werden in der Literatur ebenfalls diskutiert, treten jedoch gegenüber den ethischen Fragestellungen deutlich in den Hintergrund. Zwar werden personalisiertes Lernen, adaptive Systeme und neue Formen der Lehrgestaltung als Chancen beschrieben, doch wird deren Umsetzung stets an ethische Bedingungen geknüpft (Prilop et al., 2024; Zhai, 2024).

Insgesamt zeigt der Forschungsstand, dass die Integration von KI in Hochschulen untrennbar mit ethischen Reflexionsprozessen verbunden ist. Die zentrale Herausforderung besteht darin, technologische Innovation mit normativen Anforderungen wie Fairness, Transparenz, Datenschutz und Bildungsgerechtigkeit in Einklang zu bringen. Hochschulen stehen damit vor der Aufgabe, nicht nur KI zu implementieren, sondern auch Räume für kritische Auseinandersetzung und verantwortungsvolle Nutzung zu schaffen.

3 Ethische Prinzipien im Kontext von KI

Aus ethischer Perspektive wirft die Anwendung von KI-Systemen zahlreiche normativ relevante Fragen auf, insbesondere danach, wie ihr Einsatz gerecht, verantwortungsvoll und nachvollziehbar gestaltet werden kann. Vor dem Hintergrund dominanter und gesellschaftlich weit verbreiteter Leitbilder von Innovation, Effizienz und technologischem Fortschritt ist dabei die Frage nach einem „moralischen Fortschritt“ im Kontext der Entwicklung und Nutzung von KI-Systemen kritisch zu reflektieren. Denn technologischer Fortschritt ist nicht automatisch mit einem moralischen Fortschritt gleichzusetzen. Angesichts dessen ist zu bedenken, ob und inwiefern das technisch Mögliche auch normativ wünschenswert ist. In aktuellen Diskursen rund um die mit dem digitalen Wandel verbundenen technologischen Entwicklungen, wurde bisher sichtbar, dass sich durch ihre Anwendung, nicht nur ein großer Nutzen und neue Handlungsspielräume eröffnen, sondern auch komplexe Fragen der Regulierung, Gerechtigkeit und Verantwortung aufgeworfen werden (Bartneck et al., 2019; Floridi et al., 2018; Grimm et al., 2019; Grimm et al., 2024; Kirchschräger, 2024; Spiekermann, 2019).

In diesem Zusammenhang wird etwa gefragt, inwiefern durch den Einsatz von KI-Systemen individuelle, gesellschaftliche und demokratische Werte beeinflusst, transformiert oder möglicherweise gefährdet werden. Bezugnehmend auf diese Fragehorizonte sind KI-Anwendungen nicht als ein rein technisches Werkzeug zu begreifen, sondern als sozio-technische Systeme, die tief in gesellschaftliche Wert- und Normstrukturen eingebettet sind und aktiv mitgestaltet werden müssen. Eine differenzierte ethische Analyse der neuen Technologien

des digitalen Wandels eröffnet eine zentrale Perspektive zur Beurteilung des moralischen Fortschritts in diesem Kontext. Sie ermöglicht es, sowohl die emanzipatorischen Potenziale als auch die ambivalenten Herausforderungen, Risiken und langfristigen gesellschaftlichen Auswirkungen systematisch zu reflektieren und normativ einzuordnen.

In diesem Sinne bedarf es neben der Analyse struktureller und technologischer Rahmenbedingungen – etwa von algorithmischen Entscheidungsprozessen, Trainingsdaten und deren Qualität sowie Fragen der Systemtransparenz und Plattformlogiken (Litschka et al., 2024) – einer kritischen Auseinandersetzung mit der Frage, unter welchen Bedingungen und in welchen Formen ein verantwortungsvoller, fairer und nachhaltiger Einsatz von KI realisiert werden kann. Dies schließt auch zentrale Spannungsfelder ein, etwa zwischen Effizienzgewinnen und Datenschutz, zwischen Automatisierung und menschlicher Autonomie oder zwischen wirtschaftlichen Interessen und dem Gemeinwohl. Entscheidend ist, inwiefern es gelingt, technologische Entwicklungen so zu gestalten und institutionell zu rahmen, dass sie grundlegende Werte wie Gerechtigkeit, Autonomie, Transparenz und Verantwortung fördern.

Im Kontext aktueller Diskurse zur digitalen Ethik bzw. KI-Ethik lassen sich die darin verhandelten Problemfelder zentralen ethischen Prinzipien zuordnen: das Prinzip der Gerechtigkeit (1), das Prinzip der Nicht-Nachteiligkeit (2), das Prinzip der Vorteilhaftigkeit (3), das Prinzip der Nachvollziehbarkeit (4), das Prinzip der Autonomie (5) sowie das Prinzip des Rechts auf Privatsphäre (6) (Bartneck et al., 2019; Floridi et al., 2018; Grimm et al., 2019; Grimm et al., 2024; Kirchschräger, 2024; Spiekermann, 2019).

Das Prinzip der Gerechtigkeit (1) adressiert insbesondere Fragen algorithmischer Fairness (z. B. training bias) und zielt auf die Vermeidung der Reproduktion gesellschaftlicher Vorurteile, Diskriminierung (z. B. gender bias) und Rassismus sowie auf die Gewährleistung von Chancengleichheit beim Zugang zu und in der Nutzung von KI-Systemen.

Das Prinzip der Nicht-Nachteiligkeit (2) besagt, dass KI-Anwendungen dem Menschen (und auch Tieren oder Eigentum) keinen Schaden zufügen dürfen. Es verpflichtet dazu, potenzielle Risiken und negative Folgen – etwa durch fehlerhafte Entscheidungen von KI-Systemen oder deren missbräuchliche Nutzung – systematisch zu minimieren.

Eng damit verbunden fordert das Prinzip der Vorteilhaftigkeit (3), dass der Einsatz von KI-Systemen einen tatsächlichen Nutzen stiftet, dem Menschen „Gutes tun soll“ und zur Verbesserung bestehender Prozesse beiträgt. Im Rahmen einer Güterabwägung gilt es sicherzustellen, dass der Nutzen von KI-Anwendungen gegenüber den potenziellen Schäden überwiegt.

Das Prinzip der Erklärbarkeit (4) hebt die Bedeutung von Verantwortlichkeit und Verständlichkeit hervor. Es zielt zum einen auf die Nachvollziehbarkeit der Funktionsweise von KI-Systemen und algorithmischen Entscheidungen ab und verfolgt zum anderen das Ziel, die Vertrauenswürdigkeit von KI-Systemen zu fördern, die sich durch Kriterien wie Zuverlässigkeit, Sicherheit und Rechenschaftspflicht auszeichnet, sodass eine klare Zuweisung von Verantwortlichkeiten gewährleistet werden kann.

Das Prinzip der Autonomie (5) betont die Wahrung der Selbstbestimmung von Individuen und richtet sich gegen Formen der Bevormundung bzw. subtile Verhaltenssteuerung durch KI-

Anwendungen. Zudem darf ihre Nutzung den Menschen nicht dabei unterstützen, illegale oder unmoralische Ziele zu verfolgen und entsprechende Handlungen zu setzen. Im Zuge dessen werden Fragen des Datenschutzes und der informationellen Selbstbestimmung diskutiert, vor allem im Hinblick auf die Sammlung und Verarbeitung sensibler personenbezogener Daten, woran das Prinzip der Privatsphäre (6) geknüpft ist (Brink et al. 2024).

Zudem gerieten auch die weitreichenden demokratierelevanten und ökologischen Auswirkungen von KI-Systemen in den Fokus aktueller ethischer und politischer Diskurse. Dazu zählen mitunter die potenziellen Risiken für demokratische Prozesse, etwa durch Desinformation und automatisierte Meinungsbeeinflussung, die liberale und demokratische Grundwerte unterminieren. Zudem werden die ökologischen Auswirkungen, die mit der Nutzung leistungstarker KI-Systeme einhergehen, einer kritischen Reflexion unterzogen. Im Zentrum steht dabei insbesondere ihr erheblicher Ressourcenbedarf und Energieverbrauch, die sowohl für den Aufbau als auch den Betrieb der zugrunde liegenden technischen Infrastrukturen benötigt werden. Die energieintensive Berechnung komplexer Modelle, der steigende Bedarf an Rechenzentren sowie der damit verbundene Einsatz von Rohstoffen und Kühltechnologien werfen grundlegende Fragen hinsichtlich ökologischer Nachhaltigkeit und langfristiger Umweltverträglichkeit auf (Crawford 2021).

Angesichts dieses vielschichtigen Problemhorizonts wird deutlich, dass der Ausbau und die Anwendung von KI-Systemen in hohem Maße einer kontinuierlichen ethischen Reflexion bedürfen. Diese Reflexion ist nicht isoliert zu leisten, sondern erfordert eine enge interdisziplinäre Zusammenarbeit zwischen Technik, Ethik und Recht, um klare normative Leitlinien und geeignete regulatorische und institutionelle Rahmenbedingungen zu etablieren. Solche Maßnahmen zielen darauf ab, komplexe technologische Zusammenhänge sowie die damit einhergehenden gesellschaftlichen und normativen Problemhorizonte systematisch zu erschließen, Chancen und Risiken kritisch zu analysieren und darauf basierend einen prinzipiengeleiteten ethischen Rahmen für die Anwendung von KI-Systemen sowie entsprechende Handlungsempfehlungen zu entwickeln, die eine verantwortungsvolle Gestaltung und Nutzung dieser Technologien nachhaltig gewährleisten (Floridi et al., 2018).

4 Forschungsdesign

Der Beitrag geht der Frage nach, welche Einstellungen Hochschullehrende an einer Pädagogischen Hochschule zu generativer KI und zu ausgewählten ethischen Problemfeldern aufweisen und wie sich die aggregierten Antworttendenzen zwischen zwei quantitativen Erhebungen in den Jahren 2024 und 2025 unterscheiden.

Daraus ergeben sich folgende Teilfragen:

- Welche Chancen, Befürchtungen und ethischen Problemwahrnehmungen äußern Hochschullehrende im Hinblick auf generative KI?

- Welche Unterschiede zeigen sich in den aggregierten Antworttendenzen zwischen den Erhebungszeitpunkten 2024 und 2025?
- Welche hochschulischen Informations- und Fortbildungsangebote zum Thema KI wurden von den Befragten wahrgenommen oder besucht?

Dem Gesamtprojekt liegt ein Mixed-Methods-Design zugrunde. Der vorliegende Beitrag beschränkt sich auf den quantitativen Teil und wertet zwei Online-Erhebungen deskriptiv vergleichend aus, die im September 2024 und im September 2025 durchgeführt wurden. Kern des Projekts ist ein im Studienjahr 2024/25 durchgeführter Themenschwerpunkt, der gezielte Interventionen umfasst. Diese sollen Hochschullehrende dazu anregen, sich vertieft mit KI sowie medienethischen Fragestellungen auseinanderzusetzen.

Die Interventionen werden durch eine mehrstufige empirische Erhebung begleitet. Zu Beginn des Studienjahres (September 2024) wurde eine quantitative Online-Befragung zum Thema generative Künstliche Intelligenz (genKI) durchgeführt. Ergänzend dazu fanden im Zeitraum von Dezember 2024 bis Februar 2025 qualitative, leitfadengestützte Interviews mit Hochschullehrenden statt. Am Ende des Studienjahres wurde die quantitative Erhebung erneut durchgeführt, um Unterschiede in den aggregierten Antworttendenzen zwischen den beiden Erhebungszeitpunkten sichtbar zu machen. Aussagen über kausale Wirkungen der im Studienjahr gesetzten Maßnahmen sind auf Grundlage dieses Designs nicht möglich.

Die quantitative Erhebung wurde als standardisierte Online-Befragung mittels UniPark durchgeführt. Eingeladen wurden Hochschullehrende der KPH Wien/Niederösterreich. In die Auswertung gingen ausschließlich vollständig ausgefüllte Fragebögen ein. Zum ersten Erhebungszeitpunkt im September 2024 lagen 84 vollständig ausgefüllte Fragebögen vor, zum zweiten Erhebungszeitpunkt im September 2025 66 vollständig ausgefüllte Fragebögen, was eine Rücklaufquote von 24 bzw. 19 Prozent entspricht. Bei den beiden Erhebungen handelt es sich nicht um ein personenbezogenes Panel. Es kann daher nicht überprüft werden, ob dieselben Personen an beiden Befragungen teilgenommen haben. Die Ergebnisse werden entsprechend nicht als individuelle Veränderungen, sondern als Unterschiede zwischen aggregierten Antwortverteilungen zweier Erhebungszeitpunkte interpretiert.

Der Fragebogen umfasste Items zu allgemeinem Wissen über KI, zu generativer KI sowie zu Einstellungen, Nutzungserfahrungen und ethischen Einschätzungen (Gabriel et al., 2025). Die ethisch relevanten Items wurden theoriegeleitet ausgewählt und nachträglich den im Beitrag dargestellten ethischen Prinzipien Gerechtigkeit, Nicht-Nachteiligkeit, Vorteilhaftigkeit, Nachvollziehbarkeit, Autonomie und Privatsphäre zugeordnet. Die Zustimmung wurde über eine vierstufige Skala erhoben: „stimme nicht zu“, „stimme wenig zu“, „stimme eher zu“ und „stimme zu“. Für die Ergebnisdarstellung wurden entweder die beiden zustimmenden Kategorien oder die wenig bis nicht zustimmenden Kategorien zusammengefasst. Die Auswertung erfolgte deskriptiv anhand absoluter Häufigkeiten, Prozentwerten und Differenzen in Prozentpunkten. Ergänzend wurden Angaben zur Teilnahme an KI-bezogenen Fortbildungen erhoben, um

die Zusammensetzung der Stichprobe im Hinblick auf KI-bezogene Vorerfahrungen beschreiben zu können.

Die Online-Erhebung gliederte sich in drei zentrale Bereiche: (1) allgemeines Wissen zu KI, (2) Wissen zu generativer KI (genKI) sowie (3) vertiefende Einschätzungen zu Wissen, ethischen Fragestellungen und persönlichen Erfahrungen im Umgang mit genKI. Während in den ersten beiden Bereichen grundlegende Kenntnisse abgefragt wurden, fokussierte der dritte Teil auf Selbsteinschätzungen der Befragten. Dabei wurden sowohl fachliche Kompetenzen als auch ethische Perspektiven – etwa zu Fairness, Transparenz, Verantwortung oder Datenschutz – sowie Einstellungen und Nutzungserfahrungen im Kontext von Lehre, Forschung und Alltag erhoben.

Die zweite Online-Erhebung fand im September 2025 mit demselben Fragebogen statt, lediglich ergänzt um einen Abschnitt, der erheben sollte, welche KPH-Veranstaltungen zum Thema im vergangenen Studienjahr besucht wurden. An dieser Befragung nahmen 66 Hochschullehrende teil, auch hier waren beinahe zwei Drittel der Befragten weiblich.

Das Gesamtprojekt umfasst neben den quantitativen Erhebungen auch qualitative Interviews mit Hochschullehrenden. Diese sind jedoch nicht Gegenstand des vorliegenden Beitrags und werden daher im Folgenden nicht ausgewertet.

5 Ausgewählte Ergebnisse

5.1 Wahrnehmung hochschulischer Angebote zum Thema KI

Der Fragebogen enthielt Items zur Wahrnehmung und Nutzung hochschulischer Informations- und Fortbildungsangebote im Rahmen des KI:M:E-Schwerpunktjahres.

Zu diesen Angeboten zählten hochschulinterne Fortbildungen, das Format „Dilemma-Talk“ sowie ein Studientag im Februar 2025.

Zum zweiten Befragungszeitpunkt (September 2025) wurde erhoben, ob die Veranstaltungen Dilemma-Talk im November 2024, Studientag im Februar 2025 und Dilemma-Talk im Juni 2025 wahrgenommen bzw. besucht wurden. Der Studientag wurde von 78,79 % der Befragten besucht, wobei es sich hier um eine Pflichtveranstaltung für Stammpersonal handelt, die restlichen Befragten haben von der Veranstaltung gehört. Der Dilemma-Talk im November 2024 wurde von 22,73 % der Befragten besucht, der Dilemma-Talk im Juni 2025 noch von 13,64 %. Auffallend ist, dass bei beiden Veranstaltungen rund 20 % der Befragten diese gar nicht wahrgenommen haben. Immerhin haben rund 50 % davon gehört.

Bei den angebotenen Informationskanälen – Taskcards (Online-Pinnwand) zum Dilemma-Talk, Taskcards zum Studientag, KI-Handreichung und KI:M:E Newsletter – wurde die KI-Handreichung am stärksten wahrgenommen und gelesen. Auch der Newsletter wurde vom Großteil der Befragten gelesen. Die Taskcards sind weniger bekannt, waren allerdings stark an die genannten Veranstaltungen geknüpft und daher auch nur in diesem Kontext zugänglich.

Dagegen wurden Newsletter und KI-Handreichung direkt in die einzelnen Postfächer der Hochschullehrpersonen gemailt.

Bei den angebotenen hochschulinternen Fortbildungen geben rund 50 % der Befragten an, dass sie davon gehört haben, einige haben auch teilgenommen.

Während zum ersten Befragungszeitpunkt im September 2024 46,97 % der Befragten angeben, dass sie keine Vorträge oder Fortbildungen zu KI in den 12 Monaten davor besucht haben abseits der KPH-Angebote, geben im September 2025 68,18 % der Befragten an, dass sie vorhaben, im Studienjahr 25/26 Fortbildungen zum Thema KI zu besuchen.

5.2 Ausgewählte Ergebnisse in Bezug auf ethische Fragestellungen

In diesem Abschnitt wird ein vergleichender Einblick in ausgewählte Ergebnisse der im September 2024 (84 Fragebögen) und im September 2025 (66 Fragebögen) durchgeführten Online-Erhebungen unter Hochschullehrenden an der KPH Wien/Niederösterreich gegeben. Die in Kapitel 3 skizzierten ethischen Prinzipien werden dabei als Analyserahmen verwendet.

Tabelle (Abbildung 1) zeigt nun die ethisch relevanten Items der beiden Online-Erhebungen und ordnet sie systematisch den ethischen Prinzipien Gerechtigkeit (1), Nicht-Nachteiligkeit (2), Vorteilhaftigkeit (3), Nachvollziehbarkeit (4), Autonomie (5) sowie Privatsphäre (6) zu.

Ethische Fragestellungen (2024 & 2025)	Ethische Prinzipien
– Der Einsatz von KI trägt zu mehr Bildungsgerechtigkeit bei. – KI verstärkt Stereotype.	(1) Prinzip der Gerechtigkeit
– KI kann vorsätzlich täuschen. – KI-Systeme fügen Menschen Schaden zu. – Ein mangelndes Verständnis von KI-Systemen stellt eine Gefahr für die Demokratie dar.	(2) Prinzip der Nicht-Nachteiligkeit inkl. demokratisches Grundprinzip
– KI-Systeme bewirken Gutes.	(3) Prinzip der Vorteilhaftigkeit
– Entscheidungen und Ergebnisse von KI-Systemen müssen für Anwender*innen nachvollziehbar sein. – KI-Systeme sind vertrauenswürdig.	(4) Prinzip der Nachvollziehbarkeit
– KI kann Menschen kontrollieren.	(5) Prinzip der Autonomie
– KI-Systeme sammeln Daten ohne Zustimmung der Anwender*innen. – KI-Systeme gefährden die Privatsphäre von Personen.	(6) Prinzip der Privatsphäre

Abbildung 1: Items zu ethischen Fragestellungen (2024/2025) und Zuordnung zu ethischen Prinzipien

An dieser Stelle sei angemerkt, dass sich die ethischen Prinzipien zwar analytisch voneinander trennen und – im Sinne einer Systematisierung – den angeführten Items zuordnen lassen, sie jedoch eng miteinander verwoben sind. In diesem Sinne repräsentieren die Prinzipien kein additives, sondern ein relationales System, das eine integrative Betrachtung erfordert und Interdependenzen berücksichtigt.

So besteht etwa eine strukturelle Verschränkung zwischen den Prinzipien Gerechtigkeit (1) und Nicht-Nachteiligkeit (2): Maßnahmen, die Schaden vermeiden sollen (Nicht-Nachteiligkeit), tragen unmittelbar zur gerechten Verteilung von Risiken und Belastungen bei (Gerechtigkeit). Umgekehrt kann eine Verletzung von Fairnessprinzipien systematisch zu Benachteiligungen führen. Zudem ist das Prinzip der Vorteilhaftigkeit (3) eng an diese beiden Prinzipien geknüpft. Die Forderung, Nutzen zu maximieren, ist nur dann ethisch tragfähig, wenn sie z.B. nicht auf Kosten anderer erfolgt (Gerechtigkeit) und keine vermeidbaren Schäden verursacht (Nicht-Nachteiligkeit). In der bioethischen und technikethischen Literatur wird dies häufig als Spannungsfeld zwischen Nutzenmaximierung und Schadensvermeidung beschrieben (Werner, 2021).

Des Weiteren fungiert das Prinzip der Nachvollziehbarkeit (4) als epistemische Voraussetzung für die Anwendung der anderen Prinzipien. Ohne Transparenz und Erklärbarkeit von Entscheidungen oder Systemen ist weder überprüfbar, ob Gerechtigkeit gewahrt wird, noch ob Schaden vermieden oder Nutzen erzeugt wird. Nachvollziehbarkeit bildet somit die Grundlage für Rechenschaftspflicht sowie die Zuweisung von Verantwortung und ermöglicht dadurch eine fundierte ethische Bewertung.

Zudem steht das Prinzip der Autonomie (5) in einem dynamischen Verhältnis zu den Prinzipien der Vorteilhaftigkeit, Nicht-Nachteiligkeit und zur Privatsphäre (6). Die Achtung der Selbstbestimmung kann mit Interessen kollidieren (z.B. Sammlung von Daten ohne Zustimmung), die das Wohl des einzelnen und seine Privatsphäre gefährden und demnach Schaden verursachen. Wissenschaftlich wird dieses Spannungsfeld häufig im Kontext informierter Zustimmung und Entscheidungsfreiheit diskutiert.

Wie in Kapitel 4 zum Forschungsdesign bereits dargelegt wurde, handelt es sich bei den Erhebungen um kein personenbezogenes Panel. Folglich werden die ausgewählten Ergebnisse zu den ethischen Einschätzungen als Unterschiede zwischen aggregierten Antwortverteilungen zweier Erhebungszeitpunkte interpretiert. Die Fragebögen sahen vor, die jeweiligen Aussagen im Sinne einer Selbsteinschätzung auf einer vierstufigen Skala zu bewerten, die angibt, inwieweit die Befragten den Aussagen zustimmen. Dabei standen die Antwortmöglichkeiten „stimme nicht zu“, „stimme wenig zu“, „stimme eher zu“ und „stimme zu“ zur Auswahl. Die Abbildungen 3 und 4 stellen die Ergebnisse der beiden Online-Erhebungen in den Jahren 2024 und 2025 in absoluten Häufigkeiten und entsprechenden Prozentpunkten dar.

	stimme zu	Prozent	stimme eher zu	Prozent	stimme wenig zu	Prozent	stimme nicht zu	Prozent	keine Angabe	Prozent
KI-Systeme bewirken Gutes.	5	5,95%	44	52,38%	23	27,38%	0	0%	12	14,29%
KI kann Menschen kontrollieren	19	22,62%	25	29,76%	21	25%	10	11,90%	9	10,71%
KI verstärkt Stereotype.	35	41,67%	30	35,71%	8	9,52%	3	3,57%	8	9,52%
Ein mangelndes Verständnis von KI-Systemen stellt eine Gefahr für die Demokratie dar.	45	53,57%	22	26,19%	5	5,95%	3	3,57%	9	10,71%
KI kann vorsätzlich täuschen.	26	30,95%	24	28,57%	10	11,90%	14	16,67%	10	11,90%
KI-Systeme sammeln Daten ohne Zustimmung der Anwender:innen.	29	34,52%	29	34,52%	14	16,67%	5	5,95%	7	8,33%
Der Einsatz von KI trägt zu mehr Bildungsgerechtigkeit bei.	2	2,38%	18	21,43%	40	47,62%	16	19,05%	8	9,52%
KI-Systeme sind vertrauenswürdig.	1	1,19%	9	10,71%	42	50%	24	28,57%	8	9,52%
KI-Systeme fügen Menschen Schaden zu.	6	7,14%	20	23,81%	38	45,24%	6	7,14%	14	16,67%
KI-Systeme gefährden die Privatsphäre von Personen.	22	26,19%	33	39,29%	20	23,81%	1	1,19%	8	9,52%
Entscheidungen und Ergebnisse von KI-Systemen müssen für Anwender:innen nachvollziehbar sein.	57	67,86%	19	22,62%	0	0%	0	0%	8	9,52%

Abbildung 2: Absolute Häufigkeiten und Prozentangaben zur Erhebung im Jahr 2024

	stimme zu	Prozent	stimme eher zu	Prozent	stimme wenig zu	Prozent	stimme nicht zu	Prozent	keine Angabe	Prozent
KI-Systeme bewirken Gutes.	7	10,61%	35	53,03%	21	31,82%	1	1,52%	2	3,03%
KI kann Menschen kontrollieren	14	21,21%	31	46,97%	16	24,24%	5	7,58%	0	0%
KI verstärkt Stereotype.	41	62,12%	20	30,30%	5	7,58%	0	0%	0	0%
Ein mangelndes Verständnis von KI-Systemen stellt eine Gefahr für die Demokratie dar.	50	75,76%	15	22,73%	0	0%	1	1,52%	0	0%
KI kann vorsätzlich täuschen.	24	36,36%	23	34,85%	11	16,67%	8	12,12%	0	0%
KI-Systeme sammeln Daten ohne Zustimmung der Anwender:innen.	38	57,58%	23	34,85%	3	4,55%	1	1,52%	1	1,52%
Der Einsatz von KI trägt zu mehr Bildungsgerechtigkeit bei.	4	6,06%	13	19,70%	36	54,55%	11	16,67%	2	3,03%
KI-Systeme sind vertrauenswürdig.	0	0%	8	12,12%	35	53,03%	20	30,30%	3	4,55%
KI-Systeme fügen Menschen Schaden zu.	8	12,12%	21	31,82%	29	43,94%	4	6,06%	4	6,06%
KI-Systeme gefährden die Privatsphäre von Personen.	23	34,85%	27	40,91%	10	15,15%	3	4,55%	3	4,55%
Entscheidungen und Ergebnisse von KI-Systemen müssen für Anwender:innen nachvollziehbar sein.	57	86,36%	6	9,09%	2	3,03%	0	0%	1	1,52%

Abbildung 3: Absolute Häufigkeiten und Prozentangaben zur Erhebung im Jahr 2025

Die nachfolgende Tabelle (Abbildung 4) zeigt ausgewählte Ergebnisse der Einstellungen zu KI-Systemen von Zustimmungswerten und deren Differenz im Vergleich der Jahre 2024 und 2025. Es wird damit kein Anspruch auf eine vollständige Darstellung und Diskussion sämtlicher Aussagen erhoben. Vielmehr konzentriert sich die nachfolgende Ergebnisdarstellung auf die in den beiden Erhebungen formulierten ethischen Fragehorizonte hinsichtlich einer vertrauenswürdigen und fairen KI, wobei insbesondere die signifikanten Unterschiede zwischen beiden Erhebungszeitpunkten herausgearbeitet werden. Für die Ergebnisdarstellung werden entweder die beiden zustimmenden Kategorien oder die wenig bis nicht zustimmenden Kategorien zusammengefasst. Wie Abbildung 4 verdeutlicht, ist zwischen den beiden Erhebungszeitpunkten ein deutlicher Unterschied in der Risikowahrnehmung von KI-Anwendungen seitens der Hochschullehrenden zu verzeichnen – mit teils sehr deutlichen Anstiegen.

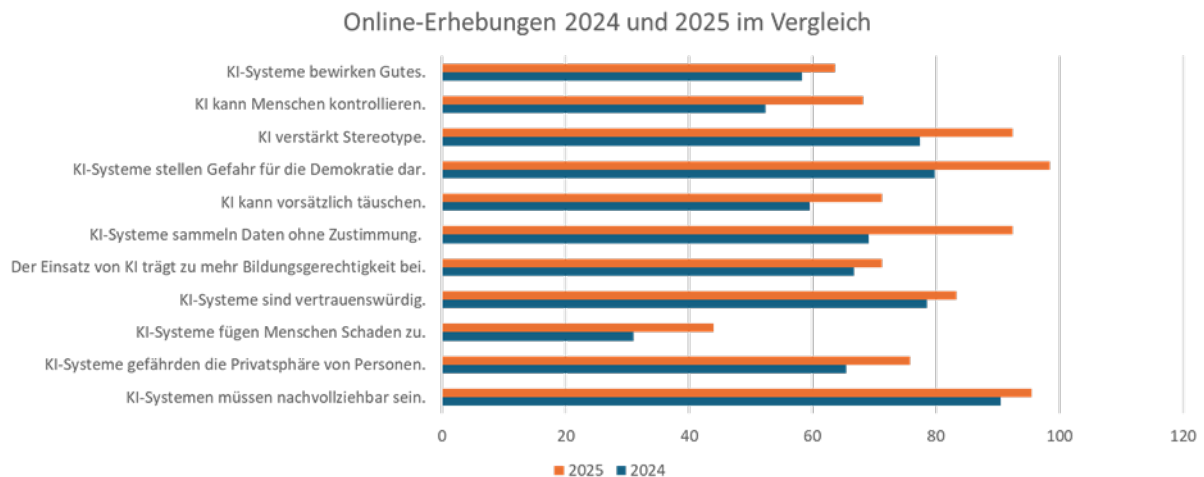


Abbildung 4: Darstellung der Unterschiede zwischen beiden Erhebungszeitpunkten 2024 und 2025

Ein Vergleich der beiden Erhebungszeitpunkte zeigt, dass sowohl im Jahr 2024 als auch im Jahr 2025 die Risikowahrnehmung von Hochschullehrenden an der KPH Wien/Niederösterreich deutlich stärker ist als die Wahrnehmung von Nutzen (Prinzip der Vorteilhaftigkeit). Während positive Effekte von KI [„Gutes bewirken“] nur moderat um 5,31 % von 58,33 % (2024) auf 63,64 % (2025) ansteigen, wird die Vertrauenswürdigkeit von KI-Systemen stark angezweifelt (2024: 78,57 % sowie 2025: 83,33 %). Dem Prinzip der Nachvollziehbarkeit wird demzufolge ein besonders hoher ethischer Stellenwert beigemessen. Zudem geben 90,48 % (2024) und 95,45 % (2025) der Befragten explizit an, dass Entscheidungen und Ergebnisse von KI-Systemen für Anwender*innen nachvollziehbar sein müssen. Wie Abbildung 2 verdeutlicht, haben die Sorgen um Kontrolle (Prinzip der Autonomie), Täuschung und Bedrohung demokratischer Werte (Prinzip der Nicht-Nachteiligkeit), Verstärkung von Stereotypen und Bildungsgerechtigkeit (Prinzip der Gerechtigkeit) sowie Gefährdung des Datenschutzes (Prinzip der Privatsphäre) besonders stark zugenommen.

Besonders auffällig sind in diesem Zusammenhang die Ergebnisse zum Datenschutz und zur Gefährdung der Privatsphäre durch KI-Anwendungen. Die Zustimmung zur Aussage, dass KI-Systeme Daten ohne Einwilligung der Anwender*innen sammeln, zeigt zwischen 2024 und 2025 einen deutlichen Anstieg von 23,39 % (von 69,04 % im Jahr 2024 auf 92,43 % im Jahr 2025). Im Vergleich dazu wird die Aussage, dass KI-Systeme die Privatsphäre von Personen gefährden, im Jahr 2024 von nur 65,48 % der Befragten bejaht und erreicht 2025 75,76 %. Die Diskrepanz zwischen den unterschiedlichen Bewertungen der beiden Aussagen lässt vermuten, dass zu beiden Erhebungszeitpunkten die Praxis der Datensammlung ohne Zustimmung von den befragten KPH-Lehrenden zwar als solche wahrgenommen wird, aber nicht notwendigerweise als Gefährdung der Privatsphäre eingestuft wird.

Ein vergleichbarer Befund zeigt sich bei der Ergebnisdarstellung des Prinzips der Nicht-Nachteiligkeit. Die Zustimmungen zu den Aussagen, dass KI-Systeme die Demokratie gefährden (a) oder vorsätzlich täuschen (b), verzeichnen bereits 2024 relativ hohe Ausgangswerte: hier sind es 79,76 % (a) und 59,52 % (b). Im Jahr 2025 steigen sie jeweils deutlich um 18,73 % (a) bzw. 11,69 % (b) an – mit Zustimmungswerten von 98,49 % (a) und 71,21 % (b). Im Gegen-

satz dazu fällt die Zustimmung zur Aussage, dass KI-Systeme Menschen Schaden zufügen, vergleichsweise gering aus. Während sie 2024 bei 30,95 % lag, ist zwar für 2025 ein Anstieg um 12,99 Prozentpunkte auf 43,94 % zu verzeichnen, das Niveau bleibt jedoch insgesamt niedrig.

Diese Unterschiede legen nahe, dass zur systematischen Analyse der wahrgenommenen Problemhorizonte eine explizite Zuordnung der jeweiligen Aussagen zu den betreffenden ethischen Prinzipien erforderlich ist. Nur durch eine solche Zuordnung lassen sich die unterschiedlichen Risikowahrnehmungen differenziert interpretieren und in den Kontext ethischer Bewertungskriterien wie Nicht-Nachteiligkeit und Privatsphäre einordnen.

Die an der KPH Wien/Niederösterreich befragten Hochschullehrenden weisen insgesamt eine kritische Haltung gegenüber KI-Systemen auf. Zudem ist zwischen beiden Erhebungszeitpunkten eine Zunahme von ethischer Sensibilität sowie ein gesteigertes Bewusstsein für potenziell negative Auswirkungen und normative Herausforderungen, die an den Einsatz von KI geknüpft sind, feststellbar.

6 Conclusio

Der Beitrag untersuchte auf Grundlage zweier quantitativer Online-Erhebungen, welche Einstellungen Hochschullehrende zu generativer KI und ausgewählten ethischen Problemfeldern äußern. Die Ergebnisse zeigen, dass ethische Fragen in beiden Erhebungen eine hohe Relevanz besitzen. Besonders hohe Zustimmung findet die Forderung nach Nachvollziehbarkeit von Entscheidungen und Ergebnissen von KI-Systemen. Ebenfalls stark ausgeprägt sind Bedenken in Bezug auf Datenschutz, demokratische Prozesse, Täuschungspotenziale und die Reproduktion von Stereotypen.

Im Vergleich der aggregierten Antwortverteilungen zeigen sich 2025 bei mehreren risiko- und ethikbezogenen Items höhere Zustimmungswerte als 2024. Diese Unterschiede können als Hinweis auf eine gestiegene Problemwahrnehmung gelesen werden. Aufgrund des Studiendesigns lassen sich daraus jedoch keine kausalen Aussagen über Wirkungen hochschulischer Maßnahmen ableiten. Weder liegt eine Kontrollgruppe vor, noch kann ohne personenbezogene Paneldaten geprüft werden, ob individuelle Einstellungsveränderungen stattgefunden haben. Trotz all dieser Einschränkungen legen die Ergebnisse nahe, dass die durchgeführten Formate Einfluss auf die Ergebnisse der zweiten Befragung hatten, denn die teilweise klaren Veränderungen im Bereich der ethischen Fragestellungen weisen auf ein verstärktes Problembewusstsein hin.

Für die Hochschulpraxis verweisen die Befunde dennoch auf einen klaren Handlungsbedarf. AI Literacy sollte nicht auf technische Anwendungskompetenz reduziert werden, sondern ethische Reflexionsfähigkeit systematisch einschließen. Hochschulen benötigen dafür transparente Leitlinien, diskursive Formate, Fortbildungsangebote und Räume für die gemeinsame Auseinandersetzung mit Fragen von Gerechtigkeit, Nicht-Nachteiligkeit, Vorteilhaftigkeit, Nachvollziehbarkeit, Autonomie und Privatsphäre.

Zukünftige Forschung sollte stärker mit längsschnittlichen Designs, Kontroll- oder Vergleichsgruppen sowie vertiefenden qualitativen Analysen arbeiten, um genauer untersuchen zu können, wie sich ethische Perspektiven von Hochschullehrenden im Umgang mit generativer KI entwickeln und welche institutionellen Angebote dazu beitragen können.

Literatur

- Bartneck, C., Lütge, C., Wagner, A. R., & Welsh, S. (2019). *Ethik in KI und Robotik*. Carl Hanser Verlag.
- Bitkom. (2024). *KI oder Kreide im Hörsaal – so digital sind Deutschlands Hochschulen*. <https://bitkom-research.de/node/941>
- Bosse, E., Wannemacher, K., & Lübcke, M. (2026). Die KI-Nutzung in Studium und Lehre. *Hochschulforum Digitalisierung*, Arbeitspapier Nr. 91.
- Brink, S., Grimm, P., Henning, C., Keber, T. O., & Zöllner, O. (2024). *Das Recht der Daten im Kontext der Digitalen Ethik*. Nomos.
- Crawford, K. (2021). *Atlas of AI: Power, politics, and the planetary costs of artificial intelligence*. Yale University Press. <https://doi.org/10.2307/j.ctv1ghv45t>
- Ceallaigh, Ó., O'Brien, T. J., Tomte, E., Kulaksiz, T., & Connolly, C. (2025). Rethinking teacher education in an AI world: Perceptions, readiness and institutional support for generative AI integration. *European Journal of Teacher Education*, 914–933. <https://doi.org/10.1080/02619768.2025.2563696>
- Estaiteyeh, M., & McQuirter, R. (2024). Generative or degenerative?! Implications of AI tools in pre-service teacher education. *Brock Education Journal*, 33(3). <https://doi.org/10.26522/brocked.v33i3.1176>
- Ferreira, A. L., Pinheiro, A., & Medeiros, P. (2025). Ethical and innovative integration of generative artificial intelligence in higher education. *International Journal of Social Science and Education Research Studies*, 5, 379–382. <https://doi.org/10.5281/zenodo.15189423>
- Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People—An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
- Gabriel, S., Jakob, J., Liebhart-Gundacker, M., & Zissler, E. (2025). KI: Medien: Ethik – Wissen. Reflektieren. Umsetzen: Awareness schaffen für medienethische Diskurse. *R&E-SOURCE*, 12(3), 327–341. <https://doi.org/10.53349/re-source.2025.i3.a1448>
- Gaisch, M., & Mader, I. (2025). *Next Generation Digital. AI Ethics & Human Factors in AI*. Excellence Edition.
- Grimm, P., Keber, T. O., & Zöllner, O. (2019). *Digitale Ethik. Leben in vernetzten Welten*. Reclam.
- Grimm, P., Trost, K. E., & Zöllner, O. (2024). *Digitale Ethik. Handbuch für Wissenschaft und Praxis*. Nomos.
- Kirchschläger, P. G. (2024). *Digitale Transformation und Ethik. Ethische Überlegungen zur Robotisierung und Automatisierung von Gesellschaft und Wirtschaft und zum Einsatz von „Künstlicher Intelligenz“*. Nomos.
- Litschka, M., Paganini, C., & Rademacher, L. (2024). *Digitalisierte Massenkommunikation und Verantwortung. Politik, Ökonomie und Ethik von Plattformen*. Nomos.

- Macgilchrist, F., Flury, C., & Roß, R. (2025). Generative KI und das Lehramtsstudium. *Erziehungswissenschaft*, 70, 27–34. <https://doi.org/10.3224/ezw.v36i1.04>
- Medeiros, R. D. O., Mascarini, A. M. N., Santos, J. P. D., Galhardi, C. M., Pozenato, C. G. D. A., Rodrigues, J. R. G., Orso, L. F., Mazzetto, F. M. C., Marin, M. J. S., Donadai, K. C. E. D. V., Barbosa, F. A. F., Silva, W. B., Roque, D. D., Herculiani, P. P., Abrão, L., & Lopes, J. G. A. (2025). Between algorithms and knowledge: Artificial intelligence in teacher education in higher education. *Interference: A Journal of Audio Culture*, 11(2), 3167–3183. <https://doi.org/10.36557/2009-3578.2025v11n2p3167-3183>
- Meylani, R. (2024). Artificial intelligence in the education of teachers: A qualitative synthesis of the cutting-edge research literature. *Journal of Computer and Education Research*, 12(24). <https://doi.org/10.18009/jcer.1477709>
- Nguyen, K. V. (2025). The Use of Generative AI Tools in Higher Education: Ethical and Pedagogical Principles. *Journal of Academic Ethics*, 23, 1435–1455. <https://doi.org/10.1007/s10805-025-09607-1>
- Petraki, E. (2024). Exploring the integration of generative AI in assessment in a tertiary context. A case study. *ASCILITE 2024 Conference Proceedings*. <https://doi.org/10.14742/apubs.2024.1537>
- Prilop, C. N., Mah, D.-K., Jacobsen, L. J., Hansen, R. R., Weber, K. E., & Hoya, F. (2024). *Generative AI in teacher education: Using AI-enhanced methods to explore teacher educators' perceptions*. OSFpreprints. <https://doi.org/10.31219/osf.io/szcwb>
- Reinmann, G., Watanabe, A., Herzberg, D., & Simon, J. (2025). Selbstbestimmtes Handeln mit KI in der Hochschule: Forschungsdefizit und –perspektiven. *Zeitschrift für Hochschulentwicklung*, 20(SH-KI-1), 33–50. <https://doi.org/10.21240/zfhe/SH-KI-1/03>
- Spiekermann, S. (2019). *Digitale Ethik: Ein Wertesystem für das 21. Jahrhundert*. Droemer.
- Suna, G., Suchismita, M., & Das, T. (2025). Integrating artificial intelligence in teacher education: A systematic analysis. *International Journal of Current Science Research and Review*, 8(1), 305–312. <https://doi.org/10.5281/zenodo.14684751>
- Tunjera, N., & Chigona, A. (2025). AI literacy in preservice teachers preparation programs: Global meta-analysis. *Proceedings of the 24th European Conference on e-Learning*, 363–370. <https://doi.org/10.34190/ecel.24.1.4174>
- Werner, M. H. (2021). *Einführung in die Ethik*. J.B. Metzler. https://doi.org/10.1007/978-3-476-05293-3_11
- Zhai, X. (2024). Transforming teachers' roles and agencies in the era of generative AI. *Journal of Science Education and Technology*, 34, 1323–1333. <https://doi.org/10.1007/s10956-024-10174-0>