



KI-Leitlinien der Privaten Pädagogischen Hochschule Burgenland

Verantwortungsvolle Nutzung von Künstlicher Intelligenz
in Studium, Lehre, Forschung und Verwaltung

(Die Leitlinien treten mit 1. Oktober 2025 in Kraft.)

INHALTSVERZEICHNIS

EXECUTIVE SUMMARY	3
1. EINLEITUNG	5
1.1 RELEVANZ UND ZIELSETZUNG	5
1.2 AUSGANGSLAGE UND OFFENE FRAGEN	5
2. RECHTLICHE GRUNDLAGEN	7
2.1 EU AI ACT	7
2.2 DATENSCHUTZ UND PERSÖNLICHKEITSRECHTE	9
3. VIER GRUNDPRINZIPIEN ZUR NUTZUNG KÜNSTLICHER INTELLIGENZ	12
3.1 AKADEMISCHE INTEGRITÄT	12
3.2 FAIRNESS	13
3.2.1 HOCHSCHULENTWICKLUNG: KONTROLLE UND MAßNAHMEN	13
3.2.2 NUTZER:INNEN: BEWUSSTSEINSBILDUNG	13
3.3 TRANSPARENZ	13
3.4 MENSCHLICHE KONTROLLE	14
4. ROLLEN UND ZUSTÄNDIGKEITEN	15
4.1 REKTORAT	15
4.2 KI-BEAUFTRAGTE AN DEN HOCHSCHULEN	15
4.3 ARBEITSGEMEINSCHAFT (AG) DIGITALE UND INFORMATISCHE BILDUNG (FORUM PRIMAR)	15
4.4 FORSCHUNGSPROJEKT FORUM PRIMAR „KI IN DER HOCHSCHULE“	16
4.5 LEHRENDE UND STUDIERENDE	16
5. INTERNE KOMMUNIKATIONSKANÄLE UND VORGEHEN BEI VERSTÖßEN	17
5.1 INSTITUTIONELLE KOMMUNIKATIONSKANÄLE	17
5.2 VORGEHEN BEI VERSTÖßEN (MISCONDUCT PROTOCOL)	17

6. KI IN LEHRE, LERNEN UND BEI PRÜFUNGEN	19
6.1 VERANTWORTUNG UND TRANSPARENZ	19
6.2 PRÜFUNGSFORMATE & BEST PRACTICES	21
6.2.1 ALTERNATIVE PRÜFUNGSFORMATE.....	21
6.2.2 BEISPIELHAFTE ZUORDNUNG VON KI-NUTZUNGSFORMEN	21
HINWEISE ZUR ANWENDUNG DER TABELLE	23
6.3 VERIFIZIERUNG AKADEMISCHER INTEGRITÄT	23
6.3.1 KI-BASIERTE PLAGIATSERKENNUNG UND KI-DETEKTOREN	23
6.3.2 ROLLE DES MENSCHLICHEN URTEILS IN PRÜFUNGSVERFAHREN.....	23
7. VERWENDUNG VON GENERATIVER KI IN ABSCHLUSSARBEITEN	25
7.1 VERANTWORTUNG VON STUDIERENDEN UND BETREUUNGSPERSONEN	25
7.2 BEWERTUNGSKRITERIEN UND SANKTIONEN	25
7.3 QUALITÄTSKONTROLLE UND ORIGINALITÄT.....	26
7.4 DIE RICHTIGE ZITATION VON KI BEI QUALIFIZIERUNGSARBEITEN	27
8. GOVERNANCE-ZYKLUS: KONTINUIERLICHE ÜBERPRÜFUNG UND ANPASSUNG	28
8.1 ZIELE UND RAHMENBEDINGUNGEN DES GOVERNANCE-ZYKLUS.....	28
8.2 PROZESSPHASEN IM GOVERNANCE-ZYKLUS.....	28
8.3 VERANTWORTLICHKEITEN UND ROLLENVERTEILUNG	29
8.4 DISSEMINATION UND KOMMUNIKATIONSSTRATEGIEN	29
8.5 QUALITÄTSSICHERUNG UND LANGFRISTIGE NACHHALTIGKEIT	29
8.6 ÜBERBLICK EXTERNER KI-LEITFÄDEN	30
9. ANHÄNGE.....	32
9.1 GLOSSAR.....	32
9.2 CHECKLISTEN UND VORLAGEN	35
LITERATURVERZEICHNIS:	40

EXECUTIVE SUMMARY

Diese Leitlinien sollen einen **Überblick** über den verantwortungsvollen Einsatz Generativer Künstlicher Intelligenz (KI) geben und richten sich an alle an der Privaten Pädagogischen Hochschule Burgenland (PPHB) tätigen und studierenden Personen und wurden auf Basis der Empfehlungen des Forschungsprojekts KI in der Hochschule des Forum Primars erstellt (Leitgeb et al., 2025a). Im Folgenden werden **zentrale Prinzipien** und Strukturen, die eine transparente, faire und datenschutzkonforme Nutzung von KI ermöglichen, skizziert.

Generative KI, insbesondere Large Language Models (LLMs), hat sich innerhalb kürzester Zeit zu einer Schlüsseltechnologie im Bildungsbereich entwickelt. Ihre Einsatzmöglichkeiten umfassen beispielsweise:

- **Vorbereitung von Lehre:** Unterstützung bei der Erstellung von Lernmaterialien, Aufgabenstellungen sowie Präsentationsunterlagen
- **Lehre:** Personalisierte Feedback-Systeme und automatisierte Auswertungen von Lernleistungen zur Unterstützung des individuellen Lernprozesses
- **Forschung:** Beschleunigung und Präzisierung der wissenschaftlichen Arbeit durch effiziente Text- und Datenanalysen
- **Verwaltung:** Vereinfachung und Optimierung administrativer Abläufe mittels automatisierter Prozesse

Damit bringt generative KI zahlreiche Möglichkeiten der Steigerung der Qualität und Effizienz im Bildungswesen mit sich. Den vielfältigen Potenzialen stehen jedoch **Risiken** gegenüber, darunter Datenschutzprobleme, algorithmische Verzerrungen (Bias, Stereotype, Diskriminierung), Fragen der akademischen Integrität und mangelnde Nachvollziehbarkeit und Transparenz bei KI-gestützten Entscheidungen.

Zentrale Prinzipien: Die Leitlinien bündeln diese wichtigsten Anforderungen an den Einsatz Generativer KI in fünf Kernbereichen:

1. Akademische Integrität

- Transparente Kennzeichnung von KI-Nutzung in Prüfungsleistungen und wissenschaftlichen Arbeiten
- Klare Unterscheidung zwischen zulässiger Unterstützung (z. B. sprachliche Korrektur) und unzulässigem Outsourcing der Eigenleistung (Ghostwriting)

2. Datenschutz

- Einhaltung der DSGVO (Datenschutz-Grundverordnung) und Privacy-by-Design-Strategien
- Sorgfältige Prüfung von KI-Tools hinsichtlich Datenweitergabe und Datenspeicherung

3. Fairness

- Implementierung von Kontroll- und Monitoring Mechanismen bzw. Etablierung von Maßnahmen (z.B. regelmäßige Klärung des rechtlichen Rahmens), um diskriminierende Effekte bei Auswahlverfahren und KI-gestützten Analysen zu verhindern

- Diversitätssensible Reflexion von Exklusions- und Benachteiligungsrisiken und aktives Entgegenwirken z.B. durch gezielte Auswahl von Modellen, Systemen und Tools

4. Transparenz

- Dokumentation und Offenlegung sämtlicher KI-basierter bzw. KI-unterstützter Prozesse und Entscheidungen, z. B. durch **KI-Logbücher** oder Protokollierung der Prompt-Eingaben, insbesondere bei Kernaufgaben der Hochschule (Lehre, Forschung, Bildungsberatung) und Entscheidungsprozessen
- Bevorzugte Nutzung von KI-Systemen, die ihre Entscheidungswege nachvollziehbar offenlegen (Explainable AI - XAI)

5. Menschliche Kontrolle

- Eine Letztentscheidung durch qualifizierte Personen bleibt unerlässlich, insbesondere bei Hochrisiko-Anwendungen (z. B. Prüfungen, Zulassungen)
- Fortlaufende Schulungen, um Lehrende, Verwaltungspersonal und Studierende für Chancen und Risiken von KI zu sensibilisieren und in der zielgerichteten und reflektierten Nutzung von KI zu professionalisieren

Governance und Rollen: Um diese Prinzipien zu verankern, werden **klare Zuständigkeiten** empfohlen:

- **Rektorat:** Hauptverantwortung für strategische Entscheidungen, Ressourcenbereitstellung und institutionelle Regelwerke.
- **AG Künstliche Intelligenz:** Bei Bedarf Koordinationsgremium für übergreifende Maßnahmen, z. B. Beratung bei der Auswahl neuer KI-Tools, Abstimmung mit Datenschutzbeauftragten.
- **KI-Beauftragte** je Hochschule: Ansprechpersonen für praktische Fragen zum KI-Einsatz (Schulungen, Implementierung, Monitoring).
- **Lehrende, Studierende und Verwaltungspersonal:** Professionalisierung und Reflexion des KI-Einsatzes, Einhaltung der Offenlegungspflichten sowie konstruktive Mitwirkung an Verbesserungsprozessen.

Kontinuierliche Aktualisierung

Aufgrund der dynamischen Entwicklung von KI (neue Versionen von LLMs, veränderte rechtliche Rahmenbedingungen wie EU AI Act) werden die Leitlinien **jährlich überprüft und an Entwicklungen angepasst**. Dieses iterative Vorgehen stellt sicher, dass aktuelle Erkenntnisse aus Forschung, Lehre und Verwaltung in die Weiterentwicklung der Richtlinien einfließen.

1. EINLEITUNG

1.1 RELEVANZ UND ZIELSETZUNG

Generative Künstliche Intelligenz (KI), insbesondere Large Language Models (LLMs) wie ChatGPT, hat sich im Rahmen globaler Digitalisierungs-, Nachhaltigkeits- und Innovationsstrategien zu einer Schlüsseltechnologie für den Hochschulsektor entwickelt. Studien belegen ihr Potenzial, **Verwaltungsprozesse zu automatisieren, personalisiertes Lernen zu ermöglichen, interdisziplinäre Forschung zu beschleunigen und nachhaltige Lösungen zu fördern** (UNESCO, 2021).

Gleichzeitig bestehen erhebliche Risiken: algorithmische Verzerrungen (Bias), Intransparenz von Entscheidungswegen, Datenschutz und Urheberrechtsfragen, hoher Energieverbrauch sowie Barrierefreiheits- und Teilhabeaspekte (WEF, 2024). Vor diesem Hintergrund sollen diese Leitlinie einen **ethisch verantwortungsvollen, nachhaltigen, barrierearmen und rechtssicheren KI-Einsatz** unterstützen. Technische, pädagogische, ökologische und rechtliche Dimensionen werden dabei gleichermaßen berücksichtigt.

Empfohlener Geltungsbereich und Ausnahmen

- **Adressat:innen:** Alle Organisationseinheiten, Forschungsgruppen, Lehrende, Forschende, Mitarbeitende der Verwaltung, Leitungspersonen, Studierende der PPH Burgenland
- **Erweiterte Nutzung:** Partnerinstitutionen und Drittmittelprojekte können den Leitfaden als Orientierungsrahmen heranziehen, sofern dies mit bestehenden Vorgaben bzw. Regelungen (z.B. des Fördergebers bzw. der Fördergeberin) vereinbar ist.
- **Mögliche Ausschlüsse:** Extern beauftragte Projekte können gesonderten Richtlinien unterliegen, sollten jedoch die vorliegenden Leitlinien einhalten.

1.2 AUSGANGSLAGE UND OFFENE FRAGEN

Aktuelle Übersichtsarbeiten (Chan, 2023; Dabis & Csáki, 2024) zeigen, dass bislang ein **ganzheitlicher Governance-Ansatz** fehlt, der Forschung, Lehre, Verwaltung und Nachhaltigkeit gleichermaßen abdeckt. Viele Hochschulen fokussieren auf Einzelthemen wie Plagiatserkennung oder Schreibassistenz, ohne ein umfassendes Regelwerk zu etablieren. Zudem sind regionale Rahmenbedingungen, etwa der EU AI Act (2024) oder nationale Datenschutzgesetze, heterogen und erschweren die Übertragbarkeit bestehender Empfehlungen.

Auch **ökologische Auswirkungen** (Stromverbrauch, CO₂Bilanz) und **Barrierefreiheit** werden in der Literatur bislang nur randständig behandelt (UNESCO, 2022). Gleiches gilt für **VendorRiskManagement**¹ von KI-Systemen und **Beschaffungskriterien** nach EU AI Act Risikoniveaus.

¹ Vendor Risk Management (VRM) bezeichnet den systematischen Prozess zur Bewertung, Steuerung und Überwachung potenzieller Risiken, die durch externe Anbieter (Vendors) entstehen können, um mögliche Auswirkungen auf die Prozesse von Hochschulen zu minimieren.

Um dieser Lücke zu begegnen, formulieren die Leitlinien vier zentrale Ziele:

1. Stärkung der akademischen Integrität – transparente Orientierungsgrundlage für den korrekten Einsatz von KI in Prüfungs- und Forschungsleistungen
2. Sicherung von Datenschutz und Urheberrecht – Privacy by Design, DSGVO-Konformität und klare Lizenzmodelle
3. Gewährleistung von Fairness, Barrierefreiheit und Nachhaltigkeit – Minimierung von Bias, inklusives Design, Energie-Monitoring und GreenAI-Strategien
4. Aufbau effektiver Governance-Strukturen – klare Zuständigkeiten, Incident Reporting, jährliche Reviews und Anpassung an neue Rechtsnormen (EU AI Act: EU, 2024)

Ein **jährlicher, datenbasierter Review-Prozess** im Rahmen des Forums Primar des PHVSO (Leitgeb et al., 2025b) prüft die Aktualität des Leitfadens. So entsteht ein **dynamisches, anpassungsfähiges Regelwerk**, das universelle Prinzipien (Transparenz, Datenschutz, Fairness, Nachhaltigkeit, Barrierefreiheit), pädagogische Zielsetzungen und regionale Besonderheiten gleichermaßen reflektiert.

Terminologiehinweis: Dieses Dokument verwendet „generative KI (GenAI)“ als Hauptbegriff. Gleichbedeutende Varianten (z. B. LLM-basierte Systeme) werden nur kontextbezogen verwendet.

2. RECHTLICHE GRUNDLAGEN

Aufbauend auf international anerkannten Richtlinien (EU, 2024; UNESCO, 2021) ist die Etablierung umfassender Standards für einen verantwortungsvollen Umgang mit generativer KI in allen hochschulischen Bereichen unabdingbar. Die rechtlichen Rahmenbedingungen sind auf europäischer Ebene mit dem AI Act (AIA) ab Sommer 2024 erstmals verbindlich geregelt (EU, 2024). Der AIA verfolgt einen risikobasierten Ansatz mit vier Risikoniveaus: inakzeptables, hohes, begrenztes und niedriges Risiko. Im Bildungssektor zählen KI-Systeme oft zur Kategorie der Hochrisiko-KI-Systeme, speziell bei Zulassungsverfahren, Lernergebnisbewertungen und Prüfungsüberwachungen.

In Österreich bestehen derzeit primär allgemeine strategische Vorgaben im Bereich der Künstlichen Intelligenz, ergänzt durch bestehende Datenschutz- (DSGVO, 2016) und Urheberrechtsregelungen (UrhG, zuletzt geändert BGBl. I Nr. 244/2021). Zusätzlich wurden Anfang 2024 zwei institutionelle Strukturen geschaffen: eine sogenannte [KI-Servicestelle](#) sowie ein [KI-Beirat](#). Diese wurden durch das KommAustria-Gesetz (§ 20c KOG) sowie ergänzend im Telekommunikationsgesetz 2021 (§ 194a TKG 2021) rechtlich verankert (BGBl. I Nr. 6/2024). Speziell für den Hochschulbereich regelt das Hochschul-Qualitätssicherungsgesetz seit 1. Juli 2024 ausdrücklich wissenschaftliches Fehlverhalten im Kontext der missbräuchlichen Nutzung von KI-Anwendungen (§ 2a Abs. 3 Z 2 HS-QSG, BGBl. I Nr. 25/2024).

Die folgenden Kapitel bündeln die wichtigsten ethischen und rechtlichen Aspekte, die als Fundament einer institutionellen Governance dienen und in fünf zentrale Handlungsfelder unterteilt werden, beginnend mit einer Zusammenfassung des EU AI Acts.

2.1 EU AI ACT

Der Artificial-Intelligence-Act (AIA), Verordnung (EU) 2024/1689, ist seit 1. August 2024 in Kraft und bildet den ersten verbindlichen Rechtsrahmen der EU für Künstliche Intelligenz. Er verfolgt einen risikobasierten Ansatz mit vier Stufen und zielt zugleich auf Innovationsförderung und Grundrechtsschutz (European Union, 2024). Für den Hochschulkontext sind vor allem Anwendungen der „Hochrisiko“-Kategorie relevant, etwa Systeme zur Zulassung, Auswahl und Leistungsbewertung oder Remote-Proctoring. Diese Verwendungen sind in Anhang III ausdrücklich erfasst (European Union, 2024; European Commission, 2025a).

Anwendung in Etappen: Verbotene Praktiken gelten seit 2. Februar 2025. Pflichten zur KI-Kompetenz (AI-Literacy) gelten seit 2. Februar 2025. Pflichten für GPAI-Modelle gelten ab 2. August 2025. Hochrisiko-Pflichten für in Anhang III genannte Bildungsanwendungen gelten ab 2. August 2026. Für Hochrisiko-Systeme, die mit sektoralen Produktvorschriften verknüpft sind, laufen Übergangsfristen bis 2. August 2027.

Der Katalog verbotener Praktiken umfasst unter anderem Systeme zur Emotionsableitung in Arbeits- und Bildungskontexten (mit engen Ausnahmen für Medizin oder Sicherheit). Hochschulen dürfen folglich keine KI-gestützte Emotionserkennung in Lehre, Prüfungen oder Bewerbungsverfahren einsetzen (European Union, 2024, Art. 5). Hochrisiko-Einsätze im Bildungsbereich sind nicht verboten, unterliegen aber strikten Auflagen (European Union, 2024; European Commission, 2025a).

Für Hochrisiko-Systeme verlangt der AIA u. a. ein fortlaufendes Risikomanagement (Art. 9), Daten- und Modell-Governance mit Dokumentation (Art. 10–11), Aufzeichnungs- und Protokollierungsmöglichkeiten (Art. 12), Nutzer:innen-verständliche Informationen und Nachvollziehbarkeit (Art. 13), wirksame menschliche Aufsicht (Art. 14) sowie Anforderungen an Genauigkeit, Robustheit und Cybersicherheit (Art. 15). Provider müssen ein Post-Market-Monitoring etablieren (Art. 72) und schwere Vorkommnisse

binnen 15 Tagen melden (Art. 73). Deploying-Stellen haben automatisch erzeugte Protokolle aufzubewahren; für Deployers schreibt Art. 26 Abs. 6 eine Mindestaufbewahrung von sechs Monaten vor (European Union, 2024). Abweichende längere oder kürzere Aufbewahrungsfristen nach Unionsrecht oder nationalem Recht bleiben unberührt.

Artikel 4 verankert KI-Kompetenz als Querschnittspflicht: Anbieter und Verwender haben Maßnahmen sicherzustellen, die das erforderliche Wissen, Können und Verständnis für den informierten und risikobewussten Einsatz von KI vermitteln. Der AIA definiert „AI literacy“ und verpflichtet zur angemessenen Qualifizierung derjenigen Personen, die KI-Systeme entwickeln oder einsetzen (European Union, 2024, Art. 4).

Spezifische Pflichten für Hochschulen als Deploying-Stellen: Vor dem Einsatz bestimmter Hochrisiko-Systeme führen öffentliche Stellen und die im AIA benannten privaten Stellen eine Grundrechts-Folgenabschätzung (FRIA) durch. Das Ergebnis wird über das von der Europäischen Kommission bereitgestellte Formular an die zuständige Marktüberwachungsbehörde übermittelt. Zudem besteht eine Registrierungspflicht im EU-Register für Hochrisiko-Systeme (Art. 49). Transparenzpflichten nach Art. 50 greifen zusätzlich, insbesondere Kennzeichnung von Deepfakes und Interaktionshinweise für KI-Systeme mit begrenztem Risiko (European Union, 2024; European Commission, 2025c).

Pflichten für GPAI-Modelle: Provider müssen technische Dokumentation führen, Downstream-Informationen bereitstellen, eine Copyright-Compliance-Policy etablieren und eine öffentlich zugängliche Zusammenfassung der Trainingsinhalte veröffentlichen; für GPAI-Modelle mit systemischem Risiko gelten verschärfte Evaluations-, Risikominderungs- und Sicherheitsauflagen (Art. 53, 55) (European Union, 2024).

Implikationen für Hochschulen (rechtlich abgeleitet und operativ): Insbesondere Hochschulen, die als Nutzer von KI-Systemen auftreten, sind von diesen Anforderungen betroffen. Daraus ergeben sich folgende konkrete Verpflichtungen:

- **Fortbildungsprogramme:** Hochschulen sind verpflichtet, kontinuierliche und zielgruppenspezifische Schulungen anzubieten, welche die technischen, ethischen und regulatorischen Aspekte von KI-Systemen abdecken (Hochschulforum Digitalisierung, 2025).
- **Dokumentation der Maßnahmen:** Alle durchgeführten Fortbildungsmaßnahmen müssen vollständig dokumentiert, nachvollziehbar archiviert und jederzeit prüfbar sein (Informationsdienst Wissenschaft, 2025).
- **Inventarisierung und Risikoklassifikation:** Hochschulen müssen systematisch erfassen, welche KI-Systeme eingesetzt werden und diese entsprechend der Risikoniveaus des AIA klassifizieren (Informationsdienst Wissenschaft, 2025).
- **Integration in das Risikomanagement:** Die dokumentierten Kompetenzen sowie Inventare sind als Indikatoren in bestehende Risikomanagement-Strukturen und Governance-Prozesse der Hochschule zu integrieren (Hochschulforum Digitalisierung, 2025).
- **Interdisziplinäre Governance-Strukturen:** Hochschulen müssen fachübergreifende Strukturen schaffen, die eine effektive Umsetzung und Einhaltung der regulatorischen Anforderungen gewährleisten (TÜV Akademie, 2025).

Diese Maßnahmen stellen sicher, dass Hochschulen ihrer Verantwortung gerecht werden, indem sie nicht nur regulatorische Compliance sicherstellen, sondern auch einen sicheren, ethisch verantwortungsvollen und informierten Umgang mit KI in Forschung, Lehre und Verwaltung fördern

2.2 DATENSCHUTZ UND PERSÖNLICHKEITSRECHTE

Die Datenschutz-Grundverordnung (DSGVO, 2016) bildet den maßgeblichen rechtlichen Rahmen für den Umgang mit personenbezogenen Daten an Hochschulen. Im Kontext der Nutzung von KI-Systemen durch Hochschullehrende ist ein verantwortungsvoller und DSGVO-konformer Umgang mit Daten unabdingbar. Falls Hochschulen KI-Systeme von externen Anbietern („Fremdsysteme“) nutzen, sollten diese vollständig DSGVO-konform sein. Die Nutzung nicht DSGVO-konformer Systeme birgt erhebliche rechtliche Risiken und kann zu Datenschutzverletzungen führen, die wiederum empfindliche Sanktionen wie Geldbußen, rechtliche Verfahren oder erhebliche Reputationsschäden zur Folge haben können. Zusätzlich besteht die Gefahr eines unrechtmäßigen Datenabflusses oder einer missbräuchlichen Verwendung der personenbezogenen Daten durch Dritte.

Im Zusammenhang mit KI-Systemen ist besondere Vorsicht geboten, da ein Hochladen bzw. Eintippen oder Hineinkopieren von Informationen bzw. Daten in insbesondere webbasierte KI-Systeme (z.B. Mistral, ChatGPT, Gemini, ...) unter Umständen mit einer Weitergabe der Daten an Dritte gleichzusetzen ist. Bei heiklen, personenbezogenen Daten ist es ratsam, diese nur mithilfe von lokal installierten bzw. auf lokalen Servern laufenden KI-Systemen intern zu verarbeiten. Um einen datenschutzkonformen Umgang zu gewährleisten und eine klare Orientierung für alle Hochschulangehörigen sicherzustellen, werden Daten wie folgt klassifiziert:

Tabelle 1: Datenklassifikation, Bedeutung und erlaubte Nutzungsmöglichkeiten

Datenklassifikation	Erlaubte Nutzungsmöglichkeiten im Hochschulkontext	Bedeutung des Begriffes
Personenbezogene Daten	- Verarbeitung nur zur Erfüllung spezifischer, rechtlich oder institutionell festgelegter Zwecke (z.B. Prüfungsverwaltung, Seminarorganisation) - Nutzung nur durch berechtigte Hochschulangehörige mit entsprechender Autorisierung - Sicherung vor unbefugtem Zugriff, nur interne Weitergabe bei dienstlicher Notwendigkeit erlaubt - Keine externe Veröffentlichung oder Weitergabe ohne explizite Zustimmung der Betroffenen	Daten, die eine eindeutige Identifizierung einzelner Personen ermöglichen, z.B. Name, Matrikelnummer, Fotos, Videos
Anonymisierte Daten	- Uneingeschränkte interne Nutzung zu Lehr- und Forschungszwecken möglich - Veröffentlichung und externe Weitergabe zulässig, wenn dauerhafte Anonymität	Daten, bei denen sämtliche identifizierenden Merkmale dauerhaft entfernt oder verändert wurden, sodass eine

	gewährleistet ist (z.B. aggregierte Evaluationsergebnisse)	Rückverfolgung zur Person unmöglich ist
Arbeitsdaten	- Interne Nutzung innerhalb der Hochschule uneingeschränkt möglich (z.B. Besprechungsprotokolle, interne Berichte) - Weitergabe außerhalb der Hochschule nur nach Freigabe durch zuständige Vorgesetzte und ohne personenbezogene Referenzen erlaubt	Daten, die bei der Erfüllung dienstlicher Aufgaben entstehen und nicht zwingend personenbezogen sind
Eigene Texte und Werke (Lehr- und Forschungsmaterialien)	- Uneingeschränkte Nutzung durch Urheber:innen für Lehrveranstaltungen, Publikationen und Weiterentwicklung - Weitergabe an Studierende und Kolleg:innen sowie Veröffentlichung gemäß dienstlichem Auftrag oder institutionellen Vorgaben gestattet	Selbstständig erstellte Materialien, wie Skripten, Übungen oder Forschungsberichte, an denen die Urheberrechte bei der Erstellerin/dem Ersteller liegen
Urheberrechtlich geschützte Werke Dritter	- Nutzung ausschließlich im Rahmen erteilter Lizenzen oder ausdrücklicher Zustimmung des Rechteinhabers gestattet (z.B. lizenzierte Lehrmaterialien, Literatur) - Bereitstellung für Studierende nur gemäß Lizenzbedingungen (z.B. Lernplattform, Semesterapparat)	Werke, die von Dritten erstellt wurden und nur mit expliziter Zustimmung oder Lizenz verwendet werden dürfen
Forschungsdaten	- Nutzung und Verarbeitung entsprechend dem genehmigten Forschungsantrag und unter Beachtung ethischer sowie datenschutzrechtlicher Standards - Speicherung und Archivierung in institutionellen Forschungsdatenbanken unter Einhaltung spezifischer Hochschulrichtlinien	Alle Daten, die im Rahmen wissenschaftlicher Forschung erhoben werden, oft sensibel und mit besonderer Beachtung ethischer und datenschutzrechtlicher Anforderungen verbunden.

	- Veröffentlichung nur anonymisiert oder nach ausdrücklicher Einwilligung der betroffenen Personen (z.B. Interviewtranskripte, Beobachtungsdaten)	
--	---	--

Beim Einsatz Generativer KI sind europäische Hochschulen verpflichtet, die DSGVO insbesondere hinsichtlich Rechtsgrundlagen, Datenminimierung, Transparenz und menschengesteuerter Überwachung automatisierter Entscheidungen umfassend zu berücksichtigen (EU, 2024). Dies umfasst:

- **Datenschutz-Folgenabschätzungen (DPIA):** Vor der Einführung neuer KI-Anwendungen sind potenzielle Risiken für personenbezogene Daten zu bewerten und Maßnahmen zur Risikominimierung vorzunehmen (Oh & Sanfilippo, 2024).
- **Privacy-by-Design-Strategien:** Hierunter fallen die strikte Umsetzung des Prinzips der Datenminimierung sowie umfassende Sicherheitsmaßnahmen gegen Datenverlust oder -abfluss, z. B. durch Verschlüsselung und Zugriffskontrollen.
- **Institutionelle Verantwortung:** Die vorliegenden KI-Leitlinien sowie regelmäßige Schulungen dienen dazu, Mitarbeitende über geltende datenschutzrechtliche Vorgaben zu informieren. Durch die Implementierung des EU AI Act sind weitere Anforderungen in den kommenden Jahren zu erwarten.

Die konsequente Einhaltung dieser Datenschutzprinzipien ist essenziell, um sowohl gesetzliche Verpflichtungen als auch ethische Anforderungen gegenüber Studierenden, Lehrenden, Verwaltungsmitarbeitenden und Forschenden sicherzustellen.

3. VIER GRUNDPRINZIPIEN ZUR NUTZUNG KÜNSTLICHER INTELLIGENZ

3.1 AKADEMISCHE INTEGRITÄT

Transparenz ist die Grundvoraussetzung akademischer Integrität. Jede Nutzung von KI – sei es für Ideenfindung, Text-, Bild- oder Code-Generierung, Datenanalyse oder Übersetzung – muss daher eindeutig gekennzeichnet werden. Um unautorisierte Hilfe zu erschweren, sind die Lehrenden angehalten, Prüfungsformate anzupassen, die die Eigenständigkeit nachweisen. Plagiatskontrollen kombinieren Software-Analysen mit stilometrischen Verfahren und persönlichen Reflexionsgesprächen. Beispielsweise wird eine Seminararbeit zunächst durch eine Software auf Textübereinstimmungen geprüft, dann auf Auffälligkeiten im Schreibstil untersucht. Bei Unklarheiten erfolgt abschließend ein persönliches Gespräch, in dem der Studierende sein Wissen zum Thema darlegen muss.

Redesign von Prüfungsformaten

Zur Vorbeugung von KI-basiertem Plagiat werden alternative Prüfungsformen empfohlen, die zugleich einen Mehrwert für den Kompetenzerwerb darstellen:

- **Mündliche Verteidigungen (z. B. Kolloquien):** Studierende erläutern ihre Arbeit im direkten Dialog, was eine tiefere Überprüfung des Lernfortschritts ermöglicht. Zudem fördert dieses Format die argumentative Kompetenz und bietet eine realitätsnahe Vorbereitung auf berufliche Kommunikationssituationen. Mündliche Prüfungen können dabei sowohl regulär vorgesehen oder bei Bedarf nach Ankündigung ergänzend durchgeführt werden.
- **Projektarbeiten:** Längere Bearbeitungsphasen und praxisnahe Aufgabenstellungen erschweren ein bloßes „Copy-and-Paste“ von KI-generierten Texten (Moorhouse et al., 2023). Diese Form fördert gleichzeitig vertieftes Verständnis, praxisbezogene Kompetenzen sowie Teamarbeit und entspricht realitätsnahen Anforderungen im späteren Berufsleben. Hierbei kann die sinnvolle Nutzung von KI besprochen und geübt werden.
- **Iterative Feedbackschleifen:** Regelmäßiges Feedback durch Lehrende oder Peers unterstützt eigenständiges Arbeiten, fördert kritische Reflexion sowie Selbstlernkompetenzen und macht unautorisierte KI-Hilfe leichter erkennbar.

Die genauen Anforderungen und Modalitäten der Prüfungsformate, insbesondere zur Bekanntgabe, erfolgen in der Regel über Lehrveranstaltungsbeschreibungen oder im Rahmen der Prüfungsordnung.

Jede Form der KI-Nutzung (z. B. für Formulierungsvorschläge, Übersetzungen, Datenauswertungen) ist verpflichtend offenzulegen. Dies erfolgt beispielsweise über Fußnoten, in einem separaten Abschnitt („KI-Logbuch“, Tool, Modellversion und den Umfang der maschinellen Unterstützung bzw. Eigenleistung dokumentiert; siehe 6.1 Verantwortung und Transparenz) oder durch klar benannte Textstellen. Damit bleibt nachvollziehbar, welcher Anteil der Arbeit vom Menschen stammt. Studierende und Lehrende sollten regelmäßig für die korrekte Kennzeichnung sensibilisiert werden.

3.2 FAIRNESS

Die Sicherstellung von Fairness und Gleichbehandlung bei der Nutzung von KI-basierten Systemen erfolgt auf zwei unterschiedlichen Ebenen:

3.2.1 HOCHSCHULENTWICKLUNG: KONTROLLE UND MAßNAHMEN

Bias-Erkennung

Die Gefahr algorithmischer Verzerrungen (Bias) entsteht häufig durch ungeeignete Trainingsdatensätze oder einseitige Modellierung von LLMs. Um Diskriminierungen zu minimieren, sollten Hochschulen regelmäßige Überprüfungen der verwendeten und angekauften KI-Systeme durchführen (Mirza et al., 2025). Diese Maßnahmen umfassen insbesondere:

- Durchführung von Stichprobenkontrollen in sensiblen Anwendungsbereichen
- Dokumentation auffälliger Abweichungen und Erarbeitung passender Gegenmaßnahmen

Inklusives Design

Bei der Beschaffung neuer KI-Systeme sollten Hochschulen solche Systeme priorisieren, die von Beginn an unterrepräsentierte Gruppen nicht benachteiligen (Mangal & Pardos, 2024). Dazu zählen insbesondere Systeme mit:

- Mehrsprachigen oder barrierefreien Nutzungsoberflächen (z. B. Kompatibilität mit Screen-readern).
- Transparenten Kriterien für Entscheidungsalgorithmen („Explainable AI“ siehe Kapitel 9.1 Glossar).

3.2.2 NUTZER:INNEN: BEWUSSTSEINSBILDUNG

Zur Vermeidung unbewusster Verstärkung von Diskriminierungseffekten sollten Nutzer:innen regelmäßig sensibilisiert und geschult werden. Zu den Maßnahmen der Bewusstseinsbildung gehören insbesondere:

- Durchführung von Schulungen und Workshops für Lehrende und Forschende zur frühzeitigen Erkennung potenzieller Verzerrungen und dem aktiven Setzen von Gegenmaßnahmen.
- Förderung eines diversitätssensiblen, kritischen und reflektierten Umgangs mit KI-generierten Inhalten, um das Risiko unbeabsichtigter Diskriminierung zu minimieren.

Das übergeordnete Ziel ist die Gewährleistung einer nicht-diskriminierenden und inklusiven Nutzung von KI-Systemen, die unterschiedliche Lern- und Lebenssituationen angemessen berücksichtigt und Chancengleichheit sicherstellt.

3.3 TRANSPARENZ

Eine nachvollziehbare Dokumentation aller KI-gestützten Prozesse ist entscheidend, um akademische Standards und ethische Verantwortbarkeit zu gewährleisten (Chan, 2023; Leitgeb & Leitgeb, 2024).

Erklärbare KI (XAI): In hoch relevanten Bereichen (z. B. Bewertungen, Zulassungen) sollten Algorithmen aufzeigen können, welche Faktoren zu einer Entscheidung führten.

Dokumentationspflicht: Log-Dateien sichern sämtliche Eingaben und Ausgaben (z. B. genutzte Prompts, KI-generierte Ergebnisse). **Versionskontrollen** ermöglichen die Nachvollziehbarkeit unterschiedlicher KI-Modelle und ihrer Änderungen. **Verantwortlichkeiten** werden eindeutig benannt, sodass bei Bedarf menschliche Instanzen einschreiten können. Gerade bei Drittanbietern ist zu prüfen, ob ausreichende Log-

und Audit-Funktionen zur Verfügung stehen, damit Institutionen ihre eigene Accountability sicherstellen können.

Kennzeichnungspflichten nach Art. 50 AI Act

- Personen sind darauf hinzuweisen, wenn sie mit einem KI-System interagieren.
- KI-generierte oder -manipulierte Medien (Bild, Audio, Video, „Deepfakes“) sind klar als künstlich zu kennzeichnen; wo möglich maschinenlesbar.
- Bei biometrischer Kategorisierung oder Emotionserkennung bestehen zusätzliche Hinweispflichten; letztere ist in Bildungskontexten ohnehin untersagt (Art. 5 i. V. m. Art. 50).

3.4 MENSCHLICHE KONTROLLE

Obwohl das Automatisierungspotenzial KI-basierter Anwendungen stetig wächst, muss die endgültige Entscheidungshoheit stets beim Menschen verbleiben (Memarian & Doleck, 2024). Dies gilt insbesondere bei sensiblen Prozessen wie Prüfungsbewertungen oder Zulassungsentscheidungen (EU-Parlament & Rat, 2024; ISO, 2023).

Um die menschliche Kontrolle angemessen sicherzustellen, sollten KI-gestützte Prozesse klar definierten Entscheidungsstufen folgen (UNESCO, 2021):

- Stufe 1 – Informieren: KI liefert lediglich Entscheidungshilfen.
- Stufe 2 – Vorschlagen: KI empfiehlt Entscheidungen, die durch Menschen bestätigt werden müssen.
- Stufe 3 – Entscheiden mit Override: KI trifft Entscheidungen, Menschen haben jedoch die Möglichkeit, diese rückgängig zu machen.

Hochrisiko-Anwendungen dürfen keinesfalls Stufe 3 überschreiten.

Regelmäßige Fortbildungen für Lehrende, Forschende und Verwaltungskräfte sind zwingend erforderlich, um KI-Anwendungen sachgerecht beurteilen und gegebenenfalls steuernd eingreifen zu können (EU-Parlament & Rat, 2024). Diese Fortbildungen umfassen sowohl technisches Grundlagenwissen (wie Funktionsweise und Grenzen von LLMs) als auch ethische und rechtliche Rahmenbedingungen (OECD, 2025).

4. ROLLEN UND ZUSTÄNDIGKEITEN

Eine wirksame Umsetzung der KI-Leitlinien erfordert klare Strukturen und verbindliche Zuständigkeiten auf allen Ebenen. Im Folgenden werden die institutionellen Rollen, die organisatorische Implementierung sowie Kommunikations- und Sanktionsmechanismen skizziert.

4.1 REKTORAT

- **Oberste Instanz:** Das Rektorat der PPH Burgenland soll die Gesamtverantwortung für institutionsweite KI-bezogene Entscheidungen und Ressourcenkoordination übernehmen.
- **Ressourcenkoordination:** Es sorgt dafür, dass ausreichende Infrastrukturen verfügbar sind (z. B. Personalressourcen, hochschuleigene KI-Systeme, Rechenkapazitäten, Lizenzen, Fortbildungsangebote) und unterstützen einen verantwortungs- und qualitätsvollen Einsatz von KI-basierten Anwendungen.

4.2 KI-BEAUFTRAGTER DER PPH BURGENLAND

- **Installation:** Seit dem Studienjahr 2023/24 sind an den Pädagogischen Hochschulen im PH-Verbund Süd-Ost eigene KI-Beauftragte tätig.
- **Aufgabenbereich:**
 - Ansprechperson für Rektorat
 - Beratung zu konkreten KI-Tools und deren Nutzung,
 - Vermittlung von Schulungen,
 - Ansprechperson für Datenschutzfragen in Kooperation mit dem:der Datenschutzbeauftragten der Hochschule.
- KI-Beauftragte für das Studienjahr 2025/26:
 - Stefan Meller (stefan.meller@ph-burgenland.at)

4.3 ARBEITSGEMEINSCHAFT (AG) DIGITALE UND INFORMATISCHE BILDUNG (FORUM PRIMAR)

- **Auftrag:** Die 2022 initiierte Arbeitsgemeinschaft Digitale und Informatische Bildung dient der **Vernetzung** zwischen den Hochschulen im PH-Verbund Süd-Ost und berät die Rektorate.
- **Zusammensetzung:**
 - KI-Beauftragte der jeweiligen Hochschule
 - Kernteam des Forschungsprojekts „KI in der Hochschule“
 - Mitarbeiter:innen AG Digitale und informatische Bildung vom Forum Primar
 - Weiterer Austausch mit Mitglieder:innen weiterer einschlägiger Gremien (z. B. Ethikkommissionen der einzelnen Hochschulen, Organisationseinheiten für digitale Kompetenz/digital gestützte Lehre)
- **Funktion:** Koordiniert die Entwicklung und Aktualisierung der zentralen Leitlinien, beobachtet deren Einhaltung und berichtet regelmäßig über Entwicklungen und Evaluationen in diesem Themenfeld an die Hochschulleitungen (Dotan et al., 2024; Maitz et al., 2024).

4.4 FORSCHUNGSPROJEKT FORUM PRIMAR „KI IN DER HOCHSCHULE“

- **Rolle:** Dieses interdisziplinäre Projekt, entstanden aus dem PH-Verbund Süd-Ost, betreibt unter anderem einen iterativen und datengestützten Prozess, um jährlich neue Leitlinien für den Einsatz von Anwendungen künstlicher Intelligenz zu entwickeln (Leitgeb et al., 2025).

4.5 LEHRENDE UND STUDIERENDE

- **Eigenverantwortung:** Alle Studierenden erhalten zu Studienbeginn Informationen zu den Leitlinien und müssen sich zusätzlich eigenständig über gültige KI-Richtlinien informieren (Cacho, 2024) sowie eine transparente Arbeitsweise sicherstellen. Dies schließen die Offenlegung und Reflexion jeglicher KI-Nutzung in allen Studiengängen ein.
- **Innovationsförderung:** Lehrende sind ausdrücklich angehalten, experimentelle KI-gestützte Szenarien in ihren Lehrveranstaltungen zu erproben (z. B. Chatbots oder automatische Feedbacksysteme), sofern akademische Standards gewahrt bleiben und die Bewertungskriterien klar definiert sind. Gleichzeitig sollen bestehende didaktische Szenarien so adaptiert werden, dass bisherige Lerninhalte weiterhin sinnvoll bearbeitet und durch KI-basierte Methoden ergänzt, aber nicht ersetzt werden. Dadurch wird eine sinnvolle Balance zwischen Innovation und Kontinuität im Lehr- und Lernprozess sichergestellt.

5. INTERNE KOMMUNIKATIONSKANÄLE UND VORGEHEN BEI VERSTÖßEN

Zur Sicherstellung einer nachhaltigen, transparenten und wirksamen Implementierung von KI-Leitlinien an Hochschulen sind systematische und zielgruppengerechte Kommunikationskanäle sowie klare Abläufe bei festgestellten Verstößen entscheidend. Forschungsergebnisse unterstreichen, dass eine transparente Kommunikation und eine nachvollziehbare Reaktion auf Verstöße erheblich dazu beitragen, die Akzeptanz von Regeln zu erhöhen und langfristig die institutionelle Integrität zu sichern (vgl. EADTU, 2024).

5.1 INSTITUTIONELLE KOMMUNIKATIONSKANÄLE

Zur regelmäßigen und umfassenden Kommunikation innerhalb der Hochschule sind folgende Maßnahmen geeignet:

- **Regelmäßige Treffen und Konferenzen:** In institutionellen Veranstaltungen (z.B. Hochschul- und Institutskonferenzen) werden aktuelle Entwicklungen, Best Practices sowie Herausforderungen im Umgang mit KI kommuniziert. Dies fördert ein kollegiales und offenes Umfeld zur aktiven Auseinandersetzung und Weiterentwicklung der KI-Leitlinien.
- **Studierendenorientierung:** Systematische Integration der KI-Leitlinien in Einführungskurse und Seminare im ersten Studienjahr. Diese Maßnahme gewährleistet, dass Studierende frühzeitig sensibilisiert werden und grundlegende Regeln verinnerlichen.

5.2 VORGEHEN BEI VERSTÖßEN (MISCONDUCT PROTOCOL)

Ein klar strukturiertes und transparentes Vorgehen bei Verstößen erhöht die Fairness und Akzeptanz von Entscheidungen innerhalb der Hochschule. Der Prozess gliedert sich in folgende Stufen:

1. Investigation

- Verdachtsfälle, wie KI-gestützte Plagiate oder missbräuchliche Datennutzung, werden an das zuständige Institut (z.B. Satzung der einzelnen Hochschulen, Prüfungskommission oder Rektorat) gemeldet.
- Eine erste Prüfung klärt, ob hinreichende Anhaltspunkte bestehen

2. Appeal

- Betroffene Studierende oder Mitarbeitende wurden informiert und haben Anspruch auf Anhörung und Erläuterung des Sachverhalts.
- Ein unabhängiges Komitee (z.B. mit KI-Beauftragten, Datenschutzexpert:innen und Fachvertreter:innen) bewertet objektiv die Einwände und Sachlage

3. Sanktionsmatrix

- **Verwarnung:** Geringfügige Verstöße, etwa nicht deklarierte KI-generierte Textpassagen ohne Täuschungsabsicht.
- **Mittlere Sanktionen (z.B. Punktabzug oder Wiederholung der Arbeit):** Bedeutendere Verstöße, beispielsweise erhebliche nicht deklarierte Nutzung von KI-generierten Inhalten, jedoch ohne nachweisbaren Vorsatz.

- **Gravierende Sanktionen (z.B. Exmatrikulation oder Ausschluss):** Schwere Verstöße wie absichtliches Ghostwriting durch KI oder vorsätzliche Datenmanipulation.

Die Sanktionen sind verhältnismäßig und berücksichtigen stets die Schwere des Verstoßes, den Grad des Vorsatzes sowie die Kooperation der betroffenen Personen während des Prozesses. Transparenz und Konsistenz in der Anwendung dieser Maßnahmen sichern die Glaubwürdigkeit der institutionellen Standards und fördern eine nachhaltige Compliance-Kultur innerhalb der Hochschule.

6. KI IN LEHRE, LERNEN UND BEI PRÜFUNGEN

Generative Künstliche Intelligenz (KI) hat in den letzten Jahren zunehmend an Bedeutung gewonnen und beeinflusst die Hochschullehre in vielfältiger Weise. Dieses Kapitel bündelt zentrale Aspekte zum verantwortungsvollen Einsatz generativer KI, um sowohl die akademische Integrität zu wahren als auch didaktische Chancen zu nutzen. Dabei werden Fragen der Transparenz, alternativer Prüfungsformate, skalierbarer Umsetzung in großen Lehrveranstaltungen sowie technischer und menschlicher Prüfmechanismen beleuchtet. Die Freiheit der Lehre, verankert im Hochschulrecht, ermöglicht es jeder Lehrperson, KI-basierte Methoden in eigener Verantwortung einzusetzen, sofern institutionell definierte Qualitäts- und Transparenzstandards eingehalten werden. Dabei trifft jede Lehrperson eigenständig die Entscheidung, ob, wie und zu welchem Zeitpunkt KI in der Lehre eingesetzt werden darf, soll oder muss.

6.1 VERANTWORTUNG UND TRANSPARENZ

Der souveräne und ethisch korrekte Umgang mit KI-Anwendungen setzt eine klare Rollen- und Verantwortungszuschreibung voraus (KMK, 2021). Sowohl Lehrende als auch Studierende sind verpflichtet, den Einsatz von KI offenzulegen und nachvollziehbar zu dokumentieren. Auf diese Weise wird akademische Authentizität gefördert und das Vertrauen in die Prüfungsleistungen gestärkt (European Network for Academic Integrity, 2020).

Um Transparenz zu fördern, wird empfohlen, ein sogenanntes KI-Logbuch in den Prüfungsarbeiten zu führen, in dem folgende Punkte offengelegt werden:

- **Verwendete Tools:** Welche Plattformen oder Software (z. B. ChatGPT, DeepL, LanguageTool) kamen zum Einsatz?
- **Einsatzbereiche:** In welchen Schritten (z. B. Gliederungserstellung, Korrekturlesen, Übersetzung) wurde KI eingesetzt?
- **Umfang der Überarbeitung/Beschreibung der Eigenleistung:** Wie stark wurden KI-Ausgaben nachbearbeitet oder durch eigene Arbeit ergänzt?

Im Folgenden wird ein Beispiel für ein KI-Logbuch präsentiert, das sowohl verschiedene Tools als auch unterschiedliche Einsatzbereiche und den jeweiligen Grad menschlicher Überarbeitung aufzeigt. Dieses Muster soll Studierenden und Lehrenden als Orientierung dienen, wie sie KI-gestützte Arbeitsschritte im Sinne einer transparenten Dokumentation aufführen können.

Tabelle 2: Vorschlag für ein KI-Logbuch

Verwendete Tools	Einsatzbereich	Eigenleistung bzw. menschliche Kontrolle und Umfang der Überarbeitungen
Elicit.com	Recherche von Literatur	Die vorgeschlagenen Quellen wurden gesammelt und anschließend manuell gesichtet sowie inhaltlich ausgewertet. Die KI diente vorrangig der Schlagwortsuche und Strukturierung.
Open AI (ChatGPT 4.1)	Review einer selbst verfassten Seminararbeit	Die von der KI gegebenen Hinweise (z. B. zu Textlogik, Stil) wurden auf Plausibilität geprüft und nur bei fachlicher Nachvollziehbarkeit in die Seminararbeit integriert.

DeepL	Übersetzung fremdsprachiger Abstracts und kurzer Textpassagen	Rohübersetzungen wurden mehrfach gegen das Original abgeglichen, fachterminologisch angepasst und stilistisch verfeinert.
LanguageTool	Korrekturlesen und Stilprüfung	Die von LanguageTool identifizierten Grammatik- und Rechtschreibfehler wurden gesichtet und, sofern zutreffend, korrigiert. Stilistische Empfehlungen wurden teilweise übernommen, teilweise verworfen.
Google Scholar	Ergänzende Literatur- und Zitatensuche	KI-gestützte Empfehlungsliste (Suggest-Funktion) wurde auf Relevanz und wissenschaftliche Qualität geprüft. Nicht vertrauenswürdige oder kontextfremde Treffer wurden aussortiert.
Zotero (KI-gestützt)	Automatisches Anlegen von Literaturverzeichniseinträgen	Die automatisierte Metadatenerkennung wurde händisch verifiziert (z. B. Titel, Erscheinungsjahr, Herausgeber). Fehlerhafte Einträge wurden korrigiert oder entfernt.
Excel mit KI-Funktionen	Erste statistische Grobauswertung	Die generierten Diagramme und statistischen Kennwerte (z. B. Mittelwerte) wurden anschließend durch eine eigene Plausibilitätsprüfung und mittels gängiger Statistikprogramme (z. B. SPSS, R) verifiziert.
Custom KI-Skript (Python)	Vorverarbeitung von Daten (Text Mining)	Die automatisierten Ausgaben (Tokenisierung, Lemmatisierung) wurden manuell überprüft und ggf. nachbearbeitet, um Fehler zu korrigieren (z. B. falsche Worterkennung, doppelte Einträge).

Kurze Erläuterungen zur Tabelle

1. **Verwendete Tools:** Hier werden sämtliche KI-gestützten Ressourcen aufgeführt, die zur Umsetzung des Projekts bzw. der Seminar- oder Abschlussarbeit genutzt wurden. Neben Namen und Version (z. B. Mistral Le Chat, Copilot, Gemini2, Perplexity, ChatGPT 5, ...) können auch Plattformen oder spezielle Programme genannt werden.
2. **Einsatzbereich:** Dieser Abschnitt spezifiziert, in welchem Stadium des Arbeitsprozesses und zu welchem Zweck das jeweilige Tool zum Einsatz kam, etwa zur Literaturrecherche, zum Korrekturlesen oder zur Datenaufbereitung.
3. **Umfang der Überarbeitungen:** Entscheidend ist eine präzise Darstellung, wie die KI-Ergebnisse (z. B. Übersetzungen, Textvorschläge, statistische Analysen) inhaltlich und formal überarbeitet wurden. Dabei sollten Fragen wie „Wurden die KI-Vorschläge nur als Inspiration genutzt oder in Gänze übernommen?“, „Wie wurde die KI-Ausgabe von mir weiterverarbeitet?“ etc. beantwortet werden.

Vorteile des KI-Logbuchs

- **Transparenz:** Das KI-Logbuch schafft Nachvollziehbarkeit über sämtliche Schritte, bei denen KI-Technologien genutzt wurden. Dies ist insbesondere bei Prüfungsleistungen essenziell, um akademische Authentizität sicherzustellen.
- **Verantwortungsbewusster Umgang:** Durch die explizite Dokumentation wird klar ersichtlich, dass die Studierenden eigenständige Entscheidungen über Relevanz, Qualität und wissenschaftliche Tragfähigkeit der KI-Ergebnisse treffen.

- **Lernförderung:** Studierende reflektieren, in welchen Bereichen KI sinnvoll unterstützt und wo Fehlerrisiken bestehen. Diese metakognitive Auseinandersetzung fördert ein tiefergehendes Verständnis wissenschaftlichen Arbeitens.
- **Fairness und Integrität:** Lehrende können anhand des Logbuchs bewerten, ob der KI-Einsatz nur unterstützenden Charakter hatte oder ob zentrale Prüfungsanforderungen durch die KI übernommen wurden. Etwaige Verstöße gegen die Richtlinien (z. B. unzulässige Vollautomatisierung) werden sichtbar.

Dieses KI-Logbuch kann, abhängig von den individuellen Hochschulvorgaben, als fester Bestandteil der Prüfungsdokumentation oder in Form von Anhang bzw. Fußnoten in die jeweilige Arbeit integriert werden. Die dort hinterlegten Informationen ermöglichen eine transparente und faire Beurteilung der tatsächlichen Eigenleistung und eine Dokumentation der an der Hochschule verwendeten KI-Systeme.

6.2 PRÜFUNGSFORMATE & BEST PRACTICES

Der Einsatz generativer KI hat die Diskussion um Prüfungsformate neu belebt. Insbesondere in textbasierten Leistungsnachweisen kann eine unkontrollierte Nutzung von KI die Gefahr von Plagiaten und oberflächlichem „Copy-and-Paste“-Verhalten erhöhen (Moorhouse et al., 2023). Um einer unreflektierten Nutzung generativer KI in Prüfungen wirksam vorzubeugen, empfiehlt es sich, verstärkt alternative Prüfungsformate einzusetzen. Die folgenden Unterkapitel erläutern methodische Ansätze, die sowohl die Eigenständigkeit der Studierenden fördern als auch die Authentizität ihrer Lernleistungen sichern.

6.2.1 ALTERNATIVE PRÜFUNGSFORMATE

- **Mündliche Verteidigung (Kolloquium):** Studierende erläutern ihre Arbeit in einem direkten Dialog, was eine tiefgehende Überprüfung des Lernfortschritts ermöglicht.
- **Projektarbeiten:** Durch realitätsnahe Aufgaben mit längerer Bearbeitungszeit erschwert sich eine reine Texterstellung durch KI.
- **Iterative Feedbackschleifen:** Regelmäßiges Feedback durch Lehrende oder Peers fördert Eigenständigkeit und Enthüllung unerlaubter KI-Hilfe.
- **Kompetenzorientierte Praxisprüfungen:** Konkrete Anwendungsaufgaben oder praktische Demonstrationen, die komplexe Kompetenzen prüfen und eine reine KI-gestützte Bearbeitung erschweren
- **Peer-Assessments:** Gegenseitige Bewertung und Rückmeldung durch Mitstudierende fördern eine vertiefte Auseinandersetzung und erschweren das unreflektierte Verwenden von KI-generierten Inhalten.
- **Offene Prüfungsgespräche:** Prüfungssettings mit offenen, individuellen Fragestellungen ermöglichen eine authentische Einschätzung der studentischen Kompetenzen und reduzieren das Risiko KI-gestützter Standardantworten.

6.2.2 BEISPIELHAFTE ZUORDNUNG VON KI-NUTZUNGSFORMEN

Bei der folgenden Zuordnung (Tabelle 2) ist zu beachten, dass neben der grundsätzlichen Einstufung (zulässig vs. unzulässig) stets die Transparenzpflicht gilt: Jede Form der KI-Nutzung ist kenntlich zu machen, insbesondere wenn sie Einfluss auf die inhaltliche Gestaltung hat.

Tabelle 3: Beispielhafte Zuordnung von KI-Nutzungsformen

KI-Nutzung	Zulässig?	Anmerkung / Kommentar
Grammatik- & Rechtschreib-Check	Ja (zulässig)	Muss gekennzeichnet werden (Fußnote/Logbuch). Greift vor allem auf sprachlicher Ebene ein, ohne den inhaltlichen Kern zu verändern.
Ideenfindung / Brainstorming	Ja (zulässig)	Eigenständige Weiterentwicklung bleibt erforderlich. Die Studierenden müssen eigene wissenschaftliche Argumentationen aufbauen und KI-Vorschläge kritisch hinterfragen.
Strukturierung und Gliederungsvorschläge	Ja (zulässig)	KI kann helfen, eine erste Gliederung zu erstellen, jedoch muss das inhaltliche Konzept (Aufbau der Argumentation, Logik, Methodik) von den Studierenden stammen.
Übersetzung fremdsprachiger Texte	Ja (zulässig)	Dient in erster Linie der Verständigung und spart Zeit. Eine gründliche fachliche Überprüfung durch die Studierenden bleibt notwendig, um Bedeutungsverluste oder Fehlübersetzungen zu vermeiden.
Literaturrecherche / -empfehlungen	Mit Vorsicht zulässig	KI-Tools können Vorschläge für Literaturquellen geben. Diese müssen allerdings manuell verifiziert werden, da KI teils fehlerhafte oder erfundene Referenzen erzeugen kann (vgl. dazu „Halluzinationen“ in LLMs).
Automatisierte Erstellung von Zusammenfassungen	Mit Vorsicht zulässig	Eine erste, grobe Inhaltsangabe kann KI-basiert erstellt werden. Die Studierenden sollten jedoch die Richtigkeit und Vollständigkeit der Zusammenfassung eigenständig überprüfen und gegebenenfalls anpassen.
Automatisierte Datenanalyse (z. B. Statistik)	Mit Vorsicht zulässig	Nur, wenn die methodische Vorgehensweise transparent bleibt und die Interpretation der Ergebnisse eigenständig erfolgt. Eine vollständige „Black-Box“-Analyse ohne eigenes Verständnis ist kritisch.
Coding / Skripte für Analysen	Mit Vorsicht zulässig	Die Anwendung kann zulässig sein, wenn die Studierenden das Skript verstehen und eigenständig anpassen können. „Copy-and-Paste“ ohne Reflektion über den Code verletzt den Lern- und Prüfungszweck.
Vollständige Textgenerierung	Nein (unzulässig)	Kernleistung ist Sache der Studierenden. KI-Tools dürfen nicht den gesamten wissenschaftlichen Text verfassen, insbesondere nicht die Argumentation, Auswertung und Interpretation übernehmen.
Erfinden empirischer Daten	Nein (unzulässig)	Verstößt gegen die Prinzipien wissenschaftlicher Redlichkeit. Eine Falschangabe oder Erdichtung von Datensätzen kann zu gravierenden Konsequenzen führen (z. B. Ausschluss vom Studium).
Ghostwriting durch KI	Nein (unzulässig)	Insbesondere bei zentralen Prüfungsleistungen (Abschlussarbeiten, Seminararbeiten) darf die KI weder Autorenschaft noch Kernrecherche übernehmen. Das eigenständige geistige Schaffen steht im Fokus.

Kaschierung / Geheimhaltung der KI- Nutzung	Nein (unzulässig)	Jeder KI-Einsatz muss transparent gemacht werden. Die Verschleierung des KI-Anteils führt zu einer Täuschung über den wahren Urheber (vgl. § 63 Hochschulgesetze).
--	------------------------------	--

Hinweise zur Anwendung der Tabelle

1. **Transparenzpflicht:** Unabhängig von der Einstufung (zulässig oder unzulässig) gilt: Jegliche Nutzung von KI-Tools ist in geeigneter Weise (z. B. KI-Logbuch, im Literaturverzeichnis, in Fußnoten oder einem separaten Abschnitt) zu dokumentieren. Dies erlaubt eine nachvollziehbare Darstellung, in welchen Schritten KI zum Einsatz kam.
2. **Didaktische Einbettung:** Zulässiger KI-Einsatz prinzipiell sollte möglichst lernförderlich gestaltet werden. Studierende sollten z. B. reflektieren, wie ihnen die KI geholfen hat und welche Grenzen sie bei der Arbeit mit KI erkannt haben.
3. **Verantwortung für den Lernfortschritt:** Wo immer KI genutzt wird, liegt die Verantwortung für die inhaltliche Richtigkeit und die wissenschaftliche Qualität weiterhin bei den Studierenden. Die KI kann unterstützend wirken, darf aber keine „Black Box“ sein, die zentrale Arbeitsschritte unsichtbar ersetzt.
4. **Fachkulturelle Unterschiede:** Die Relevanz und Konsequenzen von KI-Nutzung können in verschiedenen Disziplinen variieren (etwa in MINT-Fächern vs. Geisteswissenschaften). Daher sollte die konkrete Umsetzung der obigen Hinweise in Absprache mit Fachschaften und Modulverantwortlichen erfolgen.

Durch eine präzise Unterscheidung zwischen zulässigen und unzulässigen Formen der KI-Nutzung wird gewährleistet, dass Studierende einerseits von sinnvollen KI-Unterstützungsleistungen profitieren können und andererseits die akademische Integrität der Prüfungsleistungen gewahrt bleibt.

6.3 VERIFIZIERUNG AKADEMISCHER INTEGRITÄT

Im Kontext der zunehmenden Verwendung generativer KI-Technologien zur Erstellung akademischer Texte kommt der Etablierung zuverlässiger und transparenter Verfahren zur Sicherstellung akademischer Integrität eine zentrale Bedeutung zu. Der folgende Abschnitt vertieft die Beschreibung gängiger Tools zur KI-gestützten Plagiatprüfung, stellt Verfahren zu stichprobenartigen Überprüfungen dar und hebt insbesondere die unverzichtbare Rolle menschlicher Beurteilung bei Verdachtsfällen hervor.

6.3.1 KI-BASIERTE PLAGIATSERKENNUNG UND KI-DETEKTOREN

Der Einsatz einer KI-Applikation eignet sich nicht, um zuverlässig festzustellen, ob ein Text mithilfe von KI erstellt wurde. Zudem sind hierbei datenschutzrechtliche sowie urheberrechtliche Aspekte zu berücksichtigen.

6.3.2 ROLLE DES MENSCHLICHEN URTEILS IN PRÜFUNGSVERFAHREN

Das menschliche Urteil ist in der akademischen Praxis unverzichtbar. Dies gilt auch für den Umgang mit Verdachtsfällen hinsichtlich unerlaubter bzw. undeklarer KI-Nutzung. Leider lässt sich auch von Menschen in der Regel nicht zweifelsfrei feststellen, ob unerlaubt KI für die Verfassung einer schriftlichen Arbeit oder in einer Aufgabenstellung verwendet wurde. Es gibt allerdings Merkmale, die überproportional bei KI generierten Texten auftreten. Diese sind laut Universität Graz (2025) zum Beispiel:

- Thematische oder faktische Inkonsistenzen im Textdokument
- Unzureichende inhaltliche Originalität

- Stilbrüche innerhalb des gesamten Textdokuments
- Häufiges Vorkommen ungewöhnlicher Wortwahlen
- Häufige sprachliche Redundanzen
- Nicht-reale Quellenangaben

Weitere mögliche Hinweise auf KI-Nutzung können sein:

- Angeführte Quellen existieren nicht (fiktive Quellenangaben),
- Inhalte der angegebenen Quellen stimmen nicht mit den tatsächlichen Quellen überein (häufig bereits am Titel der Quelle erkennbar),
- Verwendung scheinbar anspruchsvoller, aber inhaltlich bedeutungsloser Formulierungen,
- Unplausible oder unrealistische Zeitangaben, beispielsweise in der Unterrichtsplanung oder bei der Beschreibung von Abläufen,
- Fachbegriffe oder Konzepte werden falsch, irreführend oder unpassend verwendet,
- Inhaltsleere Aussagen, unnötige Wiederholungen und redundante Formulierungen ohne Mehrwert,
- Verstöße gegen Datenschutz oder Urheberrecht, etwa durch die unreflektierte Verwendung geschützter Inhalte oder sensibler Daten.

In jedem Fall empfiehlt es sich bei Verdachtsfällen nach dem Vier- oder Mehr-Augenprinzip vorzugehen und im Zweifel für die Studierenden zu entscheiden.

7. VERWENDUNG VON GENERATIVER KI IN ABSCHLUSSARBEITEN

Abschlussarbeiten auf Bachelor- und Masterniveau markieren den Höhepunkt akademischer Ausbildung an unseren Hochschulen und dienen als Nachweis eigenständiger wissenschaftlicher Kompetenz. Der Einsatz generativer KI kann hier sowohl Chancen (etwa bei der Ideengenerierung) als auch Risiken (unzulässige Automatisierung) bergen. In Ergänzung zu den in den vorangehenden Kapiteln dargelegten Grundsätzen zum KI-Einsatz fokussiert dieses Unterkapitel auf die spezifischen Herausforderungen bei Abschlussarbeiten

7.1 VERANTWORTUNG VON STUDIERENDEN UND BETREUUNGSPERSONEN

Generative KI kann in verschiedenen Stadien des Forschungs- und Schreibprozesses verwendet werden – von der ersten Themenfindung bis zur sprachlichen Feinkorrektur. Dabei gilt stets, dass Studierende den Kern der wissenschaftlichen Arbeit selbstständig erbringen.

1. Eigeninitiative und Transparenz

- Studierende müssen zu jedem Zeitpunkt sicherstellen, dass generative KI lediglich unterstützend eingesetzt wird. Ein KI-Logbuch oder ein vergleichbares Offenlegungsformat (vgl. 6.1 Verantwortung und Transparenz) dient dazu, die konkreten KI-Beiträge klar auszuweisen (Marcel & Kang, 2024).
- Dies umfasst sowohl eine Angabe der verwendeten Tools (z. B. ChatGPT, DeepL) als auch eine Erläuterung, in welchen Arbeitsschritten die KI zum Einsatz kam.

2. Methodische Reflexion

- Auch in empirischen Arbeiten sollten Studierende genau dokumentieren, welche Prompt-Varianten, Datenquellen und KI-Verfahren genutzt wurden. Dies verdeutlicht den **methodischen Eigenanteil** und schützt vor unkritischer Übernahme algorithmischer Ergebnisse. Bezüglich technischer Einzelheiten (z. B. KI-gestützte Datenanalyse, Prompt-Engineering) kann auf die in
- 6.2.2 Beispielhafte Zuordnung von KI-Nutzungsformen beschriebenen „Best Practices“ sowie die in Kapitel 5.2 Vorgehen bei Verstößen (Misconduct Protocol) dargelegten Prüfroutinen verwiesen werden.

3. Rolle der Betreuungspersonen

- Lehrende bzw. Betreuende haben sicherzustellen, dass Studierende zu Beginn der Abschlussarbeit über zulässige und unzulässige KI-Nutzungen informiert werden. In Zweifelsfällen sollten sie Einsicht in entstehende Zwischenschritte oder KI-basierte Entwürfe verlangen, um den **wissenschaftlichen Erkenntnisfortschritt** nachvollziehen zu können (Moya et al., 2023).
- Bei Verdachtsmomenten (z. B. umfangreiche KI-Generierung ohne klare Kennzeichnung) gelten die in 6.3.2 Rolle des menschlichen Urteils in Prüfungsverfahren beschriebenen Regeln.

7.2 BEWERTUNGSKRITERIEN UND SANKTIONEN

Während generative KI zur sprachlichen Unterstützung und Ideenfindung häufig als zulässig eingestuft wird, ist die vollständige oder überwiegende Automatisierung wissenschaftlicher Kernleistungen bei Abschlussarbeiten unzulässig (Moorhouse et al., 2023).

1. Klar definierte Grenzen

- **Erheblicher KI-Anteil ohne Eigenleistung:** Wenn Studierende ganze Kapitel oder komplexe Analysen vollständig mithilfe von KI generieren lassen und direkt übernehmen, liegt ein Verstoß gegen akademische Integrität vor. Dies ist insbesondere dann schwerwiegend, wenn empirische Daten (z. B. Befragungsergebnisse) „erfunden“ oder manipuliert werden.
- **Graduelle Unterscheidung:** Abweichungen vom akzeptablen Maß können unterschiedlich gewichtet werden: Ein unerlaubtes Maß an KI-Einsatz bei einer Gliederungsvorlage wiegt weniger schwer als eine vollständig KI-generierte Ergebnisauswertung.

2. Sanktionsstufen

- Die Sanktionen bei nachweislich unzulässigem KI-Gebrauch folgen einem **gestuften Modell** (vgl. Memarian & Doleck, 2024):
 - *Verwarnung oder Punktabzug* bei geringfügigen Verstößen
 - *Wiederholung der Arbeit oder Teilbereiche* bei substanziellen Mängeln
 - Aberkennung des Prüfungsversuchs oder Exmatrikulation in besonders gravierenden Fällen (z. B. Täuschungsversuch größeren Ausmaßes) Die konkrete Maßnahme sollte der Schwere des Vergehens entsprechen und institutionell klar geregelt sein (vgl.
 - 6.2.2 Beispielhafte Zuordnung von KI-Nutzungsformen).

7.3 QUALITÄTSKONTROLLE UND ORIGINALITÄT

Gerade bei größeren Abschlussarbeiten fällt es Betreuer:innen schwer, KI-basierte Plagiate unmittelbar zu erkennen. Daher sind technische und reflexive Maßnahmen notwendig, um die Originalität zu bewerten.

1. Selbstreflexion und Dokumentation

- Ein KI-Einsatz-Logbuch (vgl. Abschnitt 5.1) unterstützt Betreuer:innen dabei, den Anteil maschinell generierter Inhalte nachzuvollziehen. Zusätzlich fördert es eine Selbstreflexion der Studierenden über die Sinnhaftigkeit und Grenzen der KI-Hilfe (Leitgeb & Leitgeb, 2024).
- Lehrende können gezielt mündliche Gespräche oder Kolloquien ansetzen, um Verständnis und Eigenleistung während des Betreuungsprozesses zu verifizieren.

2. Schutz der akademischen Kultur

- Abschlussarbeiten sind Zeugnis eines eigenständigen wissenschaftlichen Denk- und Arbeitsprozesses. Betreuende sollten daher Kontrollen einführen, die frühzeitig Täuschungsversuche erkennen (z. B. regelmäßige Meilensteingespräche).
- Gleichzeitig wird durch transparente Vorgaben und angemessene Unterstützungsangebote (vgl. 6.2.1 Alternative Prüfungsformate) das Bewusstsein der Studierenden für ethische und wissenschaftliche Standards geschärft.

Abschlussarbeiten erfordern eine besonders sorgfältige Auseinandersetzung mit generativer KI, da sie den Nachweis eigenständiger wissenschaftlicher Kompetenz darstellen. Indem Studierende die KI-Nutzung transparent dokumentieren, methodisch reflektieren und Lehrende eine klare Beratung sowie Überprüfungsmechanismen sicherstellen, bleibt der Wert der Originalität gewahrt. So kann generative KI in Abschlussarbeiten verantwortungsvoll und lernförderlich eingesetzt werden, ohne die akademische Integrität zu gefährden.

7.4 DIE RICHTIGE ZITATION VON KI BEI QUALIFIZIERUNGSARBEITEN

Im Rahmen von Qualifizierungsarbeiten (z. B. Bachelor-, Master- und Doktorarbeiten) ist bei Verwendung generativer KI (wie etwa ChatGPT oder ähnliche KI-gestützte Tools) unbedingt auf eine korrekte Zitationsweise zu achten. Die Nutzung von KI-gestützten Inhalten ohne angemessene Quellenangabe stellt eine Verletzung der akademischen Integrität dar und kann als Plagiat gewertet werden. Selbstredend empfiehlt es sich, auch bei kleineren Prüfungsleistungen, beispielsweise Seminararbeiten, korrekte Zitation vorzuschreiben.

Neben dem KI-Logbuch werden bei Qualifizierungsarbeiten folgende ergänzende Unterlagen empfohlen (die genauen Anforderungen können je nach Hochschule variieren):

- Dokumentation der spezifischen KI-Abfragen
- Ausdrucke oder Screenshots relevanter Ergebnisse
- Übersicht der verwendeten Einstellungen und Modelle

Die folgenden Grundsätze sind bei der Zitation von generativer KI zu beachten:

- Alle Inhalte, welche durch KI generiert wurden, müssen klar gekennzeichnet werden. Dies umfasst sowohl direkt übernommene als auch sinngemäß wiedergegebene Passagen.
- Angemessene Quellenangabe: Als Standard für die Quellenangabe wird der APA-Stil (American Psychological Association) empfohlen.
- Angabe von technischen Details: Die Zitation von generativer KI sollte zusätzlich zur Angabe der KI-Plattform auch das verwendete Modell, die Version sowie das Datum der Abfrage enthalten. Dies gewährleistet Transparenz und Nachvollziehbarkeit.

Detaillierte Zitationsweise nach APA:

Autor oder Organisation. (Jahr). Titel oder Beschreibung [Typ des Modells]. URL (Zugriffsdatum).

- **Konkrete Beispiele:**

OpenAI. (2024). ChatGPT (Version GPT-4) [Large Language Model]. <https://chat.openai.com> (Zugriffsdatum: 15. Juni 2024).

Google DeepMind. (2023). AlphaCode (Version 1.2) [KI-Programm zur Codegenerierung]. <https://deepmind.google/alphacode> (Zugriffsdatum: 20. Mai 2024).

Anthropic. (2024). Claude 3 (Version 3.0) [Large Language Model]. <https://anthropic.com/claude> (Zugriffsdatum: 10. April 2024).

Weitere spezifische Richtlinien und detaillierte Vorgaben zur APA-Zitation finden sich in der jeweils aktuellen Version der offiziellen APA-Leitlinien unter folgendem Link: <https://apastyle.apa.org>.

8. GOVERNANCE-ZYKLUS: KONTINUIERLICHE ÜBERPRÜFUNG UND ANPASSUNG

Um die Wirksamkeit der vorliegenden KI-Leitlinien sicherzustellen und deren Weiterentwicklung evidenzbasiert zu gestalten, erfolgt eine kontinuierliche Überprüfung und Anpassung durch einen definierten Governance-Zyklus (vgl. Leitgeb et al., 2024). Dazu kommt ein spezifisch entwickelter, forschungsgeleiteter Evaluationszyklus zur Anwendung. Die Grundlage dafür bildet das Projekt „Künstliche Intelligenz im PH-Verbund Süd-Ost (KI im PHVSO)“, das von 2024 bis 2028 (mit optionaler Verlängerung bis 2032) durchgeführt wird.

8.1 ZIELE UND RAHMENBEDINGUNGEN DES GOVERNANCE-ZYKLUS

Die Ziele des Governance-Zyklus umfassen die systematische und empirisch fundierte Weiterentwicklung der KI-Leitlinien. Im Zentrum stehen dabei die Bereiche Lehren und Lernen. Der Zyklus verfolgt eine strukturierte, jährliche Routine zur Datenerhebung, Analyse, Anpassung und Dissemination, um auf Veränderungen technologischer, didaktischer und ethischer Rahmenbedingungen reagieren zu können.

8.2 PROZESSPHASEN IM GOVERNANCE-ZYKLUS

Der Governance-Zyklus gliedert sich in vier Phasen, die jährlich durchlaufen werden:

Phase 1: Empirische Datenerhebung

In der ersten Phase (Oktober bis Februar jedes Studienjahres) erfolgt eine umfassende Datenerhebung, welche quantitative und qualitative Methoden kombiniert:

- **Quantitative Befragungen:** Jährliche Vollerhebung unter Studierenden des ersten Semesters (ca. n=500) und Lehrenden zu KI-Nutzungsverhalten, Einstellungen und Bedarfen (Fragebogen nach etablierten Standards)
- **Qualitative Interviews:** Systematische Auswahl und Befragung von Hochschullehrenden zu Best-Practice-Beispielen des KI-Einsatzes in der Hochschullehre

Phase 2: Datenanalyse und Evidenzsynthese

Von März bis Juni erfolgt eine umfangreiche Analyse der erhobenen Daten:

- **Quantitative Analyse:** Durchführung statistischer Analysen wie Trendanalysen, Regressionsmodelle und Mehrebenenanalysen (Mixed-Effects-Modelle), um Veränderungen und Determinanten des Nutzungsverhaltens zu identifizieren
- **Qualitative Inhaltsanalyse:** Auswertung der Interviews zu Good-Practice-Beispielen, um exemplarische Lehr-/Lernkonzepte systematisch zu erfassen
- **Systematic Literature Review (SLR):** Parallele, jährliche systematische Übersichtsarbeit zur internationalen KI-Forschung im Bereich der Hochschullehre und Abschlussarbeiten

- Diese erfolgt nach PRISMA-Richtlinien und ergänzt die empirischen Befunde um den aktuellen Forschungsstand

Phase 3: Entwicklung eines evidenz-basierten Berichts

Ab Juli 2026 werden jedes Studienjahr die Ergebnisse aus Datenerhebungen und SLR in einem "Evidence Brief" (ca. 10 Seiten) konsolidiert. Dieser fasst zentrale Befunde, Trends, Good-Practice-Beispiele sowie Empfehlungen kompakt und nachvollziehbar zusammen.

Der "Evidence Brief" wird zur offenen Konsultation und Diskussion an die relevanten Gremien und Stakeholder (Rektorate, AG KI, KI-Beauftragte, Lehrende, Studierende) weitergegeben. Ziel ist eine breite und partizipative Validierung der Ergebnisse und der daraus abgeleiteten Handlungsempfehlungen.

Phase 4: Anpassung und Veröffentlichung der Leitlinien

Die AG KI und das Projektteam des KI-Projekts im PHVSO prüfen notwendige Anpassungen der bestehenden KI-Leitlinien. Im September jedes Jahres wird die finale, aktualisierte Version der Leitlinien beschlossen und anschließend zum 1. Oktober veröffentlicht. Jede aktualisierte Version erhält einen DOI zur besseren Auffindbarkeit und zur Sicherstellung der langfristigen Dokumentation.

8.3 VERANTWORTLICHKEITEN UND ROLLENVERTEILUNG

Die Verantwortlichkeiten im Governance-Zyklus sind klar definiert und folgen einer RACI-Logik (Responsible, Accountable, Consulted, Informed):

- **Rektorat:** Gesamtverantwortung und strategische Entscheidungen (Accountable)
- **Projektteam KI im PHVSO:** Datenerhebung, Analyse und Erstellung des Evidence Briefs (Responsible/Accountable)
- **AG KI:** Fachliche Begleitung, Konsultation und Unterstützung bei der Erstellung des Evidence Briefs und der Revision der Leitlinien (Consulted)
- **KI-Beauftragte und Lehrende:** Konsultiert für Praxisbezug, Umsetzung und Qualitätssicherung (Consulted/Responsible)
- **Studierende:** Informiert und involviert in Form von Befragungen und Feedbackschleifen (Informed)

8.4 DISSEMINATION UND KOMMUNIKATIONSSTRATEGIEN

Die Ergebnisse des Governance-Zyklus werden sowohl intern innerhalb der Hochschulen als auch extern öffentlich zugänglich gemacht. Intern erfolgt die Dissemination über die jeweiligen Plattformen der PHen (Forum Primar, interne Newsletter, Moodle). Extern werden Ergebnisse mittels Publikationen in nationalen und internationalen Fachzeitschriften (z.B. Zeitschrift für Hochschulentwicklung, European Journal of Teacher Education) und Konferenzbeiträgen (z.B. ÖFEB, ECER) verbreitet.

8.5 QUALITÄTSSICHERUNG UND LANGFRISTIGE NACHHALTIGKEIT

Durch die systematische, jährlich wiederholte Datenerhebung und SLR ist die Qualitätssicherung der KI-Leitlinien langfristig gewährleistet. Die Ergebnisse ermöglichen nicht nur eine evidenzbasierte Optimierung der Leitlinien, sondern schaffen gleichzeitig eine langfristige Wissensbasis zur KI-Nutzung im Hochschulkontext, die auch nach Projektende genutzt und weitergeführt werden kann.

8.6 ÜBERBLICK EXTERNER KI-LEITFÄDEN

Die Entwicklung und Implementierung von Leitlinien zur Nutzung generativer KI ist ein globales Thema, dem sich Hochschulen international stellen. Im Sinne eines Vergleichs und einer Orientierung an Best Practices wurden im Folgenden zentrale KI-Leitfäden renommierter internationaler Hochschulen zusammengetragen und tabellarisch aufbereitet. Die hier dargestellten internationalen Referenzbeispiele fokussieren auf zentrale Aspekte wie akademische Integrität, Transparenz, Datenschutz, didaktische Integration und menschliche Kontrolle und dienen somit als wertvolle Orientierung für die evidenzbasierte Gestaltung von KI-Richtlinien an Pädagogischen Hochschulen.

Tabelle 5: Überblick externer KI-Leitfäden

Hochschule	Kurzzusammenfassung	Zitat	Quelle
Uppsala University (SE)	Regelt Einsatz von GenAI in Lehre/Prüfung, betont Datenschutz, akademische Integrität, Transparenz, menschliche Verantwortung, kontinuierliche Schulungen.	“The guideline for the use of generative AI in teaching and assessment came into effect on 20 January 2025.”	https://www.uu.se/en/staff/gateway/teaching/ai-in-teaching-and-learning/guidelines-on-generative-ai
Harvard University (US)	Erlaubt „verantwortliches Experimentieren“, verlangt Datensicherheit, Copyright-Beachtung, Prüfung von Ausgaben, keine sensiblen Daten in offene Modelle.	“The University supports responsible experimentation with generative AI tools.”	https://huit.harvard.edu/ai/guidelines
University of Cambridge (UK)	Studierende dürfen GenAI für Selbststudium nutzen; Pflicht zur Quellenprüfung, Offenlegung, Fachrichtungs-sonderregeln, Schulung für Lehrende.	“Students are permitted to make appropriate use of GenAI tools to support their personal study.”	https://blendedlearning.cam.ac.uk/artificial-intelligence-and-education/using-generative-ai
University of Oxford (UK)	Leitlinien für Lehre/Prüfung; AI darf Lernanstrengung nicht ersetzen; fordert Zitation, Bias-Bewusstsein, Verankerung in Assessment-Design.	“AI cannot be seen as a shortcut or replacement of the individual effort needed to acquire intellectual skills.”	https://academic.admission.ox.ac.uk/ai-in-teaching-and-assessment
ETH Zürich (CH)	Richtet sich an Lehrende; proaktiver, aber kontrollierter Einsatz; Transparenz, Fairness,	“ETH Zurich advocates a proactive approach to the use of	https://ethz.ch/content/dam/ethz/main/ethz/education/ai_in_educati

**Universität
Wien (AT)**

Qualitätssicherung, Governance Board, didaktische Szenarien.	generative AI within educational contexts.”	on/Generative%20AI%20in%20Teaching%20and%20Learning%20-%20Guidelines%20ETH.pdf
--	---	---

Praxisleitfaden für Lehrende/Studierende; Fokus auf rechtl. Rahmen (EUAIAct, DSGVO), ToolKompass, Prüfungsdesign, PromptBeispiele, Reflexion.	“Guidelines der Universität Wien zum Umgang mit Künstlicher Intelligenz (KI) in der Lehre”	https://phaidra.univie.ac.at/o:2092606
---	--	---

Uni Graz

Orientierungsrahmen für Lehrende und Studierende	“Lehren und Lernen mit KI”	https://lehren-und-lernen-mit-ki.uni-graz.at/de/
--	----------------------------	---

**TU
Darmstadt**

Informationen zum Umgang mit KI-Tools im Zusammenhang mit wissenschaftlichem Arbeiten	“Empfehlung zur Kennzeichnung und Dokumentation von KI-Generaten”	https://www.ulb.tu-darmstadt.de/forschen/publizieren/ki/index.de.jsp
---	---	---

9. ANHÄNGE

Um den inhaltlichen Umfang dieser Leitlinien klar und übersichtlich zu gestalten, wird im Folgenden ein eigenes Kapitel für Glossar, Checklisten und Vorlagen bereitgestellt. Dadurch lassen sich wichtige Fachbegriffe und Dokumentationshilfen an zentraler Stelle auffinden.

9.1 GLOSSAR

Nachfolgend finden sich zentrale Begriffe in alphabetischer Reihenfolge. Die Auswahl orientiert sich an den häufig verwendeten Fachtermini in diesem Dokument. Alle Definitionen dienen als Orientierung; je nach Hochschule oder Fachdisziplin können abweichende Formulierungen gelten.

1. **Abschlussarbeiten**

Wissenschaftliche Arbeiten (z. B. Bachelor-, Master- und Diplomarbeiten), die den erfolgreichen Abschluss eines Studiengangs dokumentieren. Sie stellen den Nachweis eigenständiger wissenschaftlicher Kompetenz dar und erfordern eine besonders sorgfältige Auseinandersetzung mit generativer KI, um akademische Integrität zu gewährleisten.

2. **Academic Integrity (Akademische Integrität)**

Übergeordnete Leitprinzipien wissenschaftlichen Arbeitens, die Ehrlichkeit, Transparenz und Verantwortung umfassen. Bei der Nutzung von KI bezieht sich akademische Integrität insbesondere darauf, wie viel Unterstützung durch KI zulässig ist und wie diese offen offengelegt wird.

3. **AI-Literacy**

Die Fähigkeit, grundlegende Funktionsweisen, Möglichkeiten und Risiken von KI-Systemen zu verstehen, sie kritisch zu reflektieren und bedarfsorientiert einzusetzen. AI-Literacy umfasst neben technisch-methodischen auch ethische und rechtliche Aspekte.

4. **Algorithmic Bias (Algorithmenverzerrung)**

Systematische Verzerrungen in KI-Entscheidungen, die durch unzureichende Trainingsdatensätze, einseitige Modellierung oder andere Fehlerquellen entstehen können. Diese Verzerrungen können zu Benachteiligungen einzelner Personen oder Gruppen führen (z. B. geschlechtsspezifische oder ethnische Diskriminierung).

5. **Automatisierte Datenanalyse**

Der KI-gestützte Prozess, bei dem große Datenmengen (Big Data) mithilfe statistischer und maschineller Verfahren untersucht werden. Ziel ist es, Muster, Zusammenhänge oder Vorhersagen zu generieren, die menschlich nur schwer manuell zu ermitteln wären.

6. **Black-Box-Modell**

KI-Modell, dessen interne Entscheidungslogik selbst für Fachleute nicht oder nur schwer nachvollziehbar ist; erschwert Erklärbarkeit und Verantwortungszuschreibung.

7. **Big Data**

Umfangreiche und komplexe Datenbestände, die in hoher Geschwindigkeit anfallen (Volume, Velocity, Variety) und mit konventionellen Datenverarbeitungstechnologien nur schwer zu bewältigen sind. KI-Methoden werden häufig zur Extraktion relevanter Informationen aus Big Data eingesetzt.

8. **Data Minimization (Datenminimierung)**

Grundprinzip der DSGVO und des Privacy-by-Design-Ansatzes, nach dem nur so viele personenbezogene Daten erhoben und verarbeitet werden dürfen, wie unbedingt notwendig. Dies zielt auf eine Reduktion von Datenschutzrisiken bei der Nutzung von KI-Systemen ab.

9. **Datenschutz-Folgenabschätzung (DPIA)**

Ein strukturiertes Verfahren, um vor der Einführung neuer Technologien - insbesondere KI-Anwendungen - Risiken für personenbezogene Daten zu identifizieren und zu bewerten. Dient der Einhaltung gesetzlicher Vorschriften (z. B. DSGVO) und dem Schutz von Betroffenenrechten.

10. **Deep Learning**

Unterkategorie des maschinellen Lernens, die tiefe (mehrschichtige) neuronale Netze verwendet, um hierarchische Repräsentationen aus Rohdaten zu erlernen; liefert Spitzenleistungen u. a. in Sprach-, Bild- und Audioverarbeitung.

11. **DSGVO (Datenschutz-Grundverordnung)**

Rechtsrahmen der Europäischen Union, der Vorgaben für den Umgang mit personenbezogenen Daten enthält. Bei KI-Anwendungen sind insbesondere Aspekte wie Einwilligung, Transparenz und Datenminimierung (Privacy by Design) zu beachten.

12. **EU AI Act**

Erste umfassende EU-Regelung für Künstliche Intelligenz; arbeitet mit einer vierstufigen Risikoklassifikation (inakzeptabel – hoch – begrenzt – minimal) und legt für Hochrisiko-Systeme – darunter KI-Anwendungen im Bildungsbereich – strikte Pflichten zu Risikomanagement, Daten-Governance, Transparenz und menschlicher Aufsicht fest.

13. **Explainable AI (XAI)**

Teilgebiet der Künstlichen Intelligenz, das sich auf nachvollziehbare Entscheidungsprozesse fokussiert. Ziel ist es, komplexe Modelle (z. B. neuronale Netze, LLMs) so zu gestalten oder zu ergänzen, dass Menschen die zugrunde liegenden Entscheidungswege nachvollziehen können.

14. **Fairness**

In Bezug auf KI-Anwendungen das Prinzip, alle Personen oder Gruppen gleichwertig zu behandeln und Diskriminierung (z. B. durch algorithmische Verzerrungen) zu vermeiden. Die Fairnessprüfung erfordert oft zusätzliche Audits oder spezielle statistische Verfahren zur Bias-Erkennung.

15. **Fern-/automatisiertes Prüfungsaufsichtssystem**

Software, die mittels Video-, Audio- oder Tastenanschlaganalyse sowie bild-/stimmbasierter Authentifizierung Prüfende identifiziert und verdächtiges Verhalten in Online-Prüfungen erkennt

16. **Foundation-Modell**

Großskaliges, vortrainiertes KI-Modell – meist Transformer-Architektur –, das mit sehr heterogenen Datenquellen trainiert wird und anschließend durch Feinabstimmung / Prompting auf zahlreiche Spezialaufgaben angepasst werden kann.

17. **Generative KI**

KI-Systeme, die basierend auf umfangreichen Trainingsdaten neue Inhalte (Texte, Bilder, Sprach- oder Videosequenzen) eigenständig erzeugen können. Sie werden unter anderem in Form sogenannter **Large Language Models (LLMs)** eingesetzt, um natürlichsprachliche Texte zu generieren.

18. Ghostwriting

Unzulässige Form akademischer Unterstützung, bei der KI (oder eine andere Person) wesentliche Leistungen in Studium oder Forschung erbringt und diese als eigene Leistung ausgegeben wird. Ghostwriting durch KI wird in den Leitlinien explizit als Täuschungsversuch gewertet.

19. Green AI / Energiemonitoring

Forschung- und Entwicklungspraktiken, die den Energie- und CO₂-Fußabdruck von KI-Training und -Inference systematisch messen und senken, z. B. durch Reporting von Stromverbrauchskennzahlen.

20. High-Risk (gem. EU AI Act)

Kategorie im geplanten europäischen Rechtsrahmen für KI (EU AI Act), in der Anwendungen eingeordnet werden, die ein hohes Gefährdungspotenzial für Menschen und Gesellschaft aufweisen. Im Bildungssektor fallen häufig automatisierte Prüfungsverfahren oder Zulassungsprozesse darunter.

21. KI-Detektionstools

Software, die versucht, Unterschiede zwischen menschlich verfassten und KI-generierten Texten zu erkennen. Diese Tools sind oft noch fehleranfällig (falsch-positive bzw. falsch-negative Ergebnisse), weshalb ihr Einsatz immer durch menschliche Expertise ergänzt werden sollte.

22. KI-Logbuch

Dokumentationsformat, in dem sämtliche KI-basierten Arbeitsschritte festgehalten werden. Ein solches Logbuch kann Auskunft geben über (a) verwendete Tools, (b) Einsatzbereiche und (c) die nachträgliche Überarbeitung. Trägt wesentlich zur Transparenz in Prüfungen und wissenschaftlichen Arbeiten bei.

23. KI-Beirat / KI-Kommission

Gremium innerhalb einer Institution, das Richtlinien und strategische Entscheidungen zum Einsatz von KI koordiniert. In Hochschulen oft zusammengesetzt aus Vertreter:innen der Hochschulleitung, Datenschutzbeauftragten, Fachdisziplinen und studentischen Vertretungen.

24. KI-Nutzungsoffenlegung (AI Use Disclosure)

Dokumentierte Angabe darüber, wo, wie und in welchem Umfang KI bei der Erstellung eines Produkts, einer Dienstleistung oder wissenschaftlichen Arbeit eingesetzt wurde erhöht Transparenz und Rechenschaftspflicht.

25. Large Language Model (LLM)

Spezieller Typ maschinellen Lernens, der riesige Textkorpora nutzt, um Muster in Sprache zu erkennen und selbstständig qualitativ hochwertige Texte zu generieren. Bekannte Beispiele sind GPT-Modelle (z. B. GPT4.x) oder BERT-Varianten.

26. Maschinelles Lernen

Teilgebiet der KI, in dem Algorithmen ihre Leistung durch das Analysieren von Beispieldaten verbessern, ohne für jede Einzelfunktion explizit programmiert zu sein.

27. Monitoring

Laufende Überwachung von KI-Systemen, um sowohl technische Fehler (z. B. Abstürze, Modellfehler) als auch regulatorische oder ethische Verstöße (z. B. Datenschutzverletzungen) frühzeitig zu erkennen. Monitoring wird insbesondere bei Hochrisikooanwendungen gefordert.

28. Neuronales Netz

Rechnerisches Modell aus Schichten künstlicher „Neuronen“, deren gewichtete Verknüpfungen Eingaben iterativ in Ausgaben transformieren; Grundbaustein moderner Deep-Learning-Systeme.

29. Plagiatserkennungstool (KI-gestützt)

Software, die mittels maschineller Lern- bzw. Natural Language Processing-Verfahren kopierte, paraphrasierte oder KI-generierte Textpassagen identifiziert, indem sprachliche, semantische und stilometrische Merkmale analysiert werden.

30. Privacy by Design

Konzeption, nach der Datenschutzaspekte von Anfang an in die Systementwicklung integriert werden. Bei KI-Tools bedeutet das, dass Datenströme, Speicherpraktiken und Algorithmen so gestaltet sind, dass Nutzerdaten bestmöglich geschützt werden.

31. Prompt Engineering

Prozess des gezielten Formulierens von Eingabeaufforderungen (Prompts) an KI-Systeme, um möglichst präzise und verwertbare Ausgaben zu erhalten. Essenziell für Lehrende und Forschende, die generative KI in didaktischen Szenarien oder wissenschaftlichen Analysen einsetzen wollen.

32. Risikokategorisierung

Einteilung von KI-Anwendungen in unterschiedliche Gefährdungsstufen (z. B. „High Risk“ oder „Low Risk“) auf Basis potenzieller Schäden für Individuen oder gesellschaftliche Systeme. Im EU AI Act wird diese Kategorisierung verbindlich vorgegeben.

33. Sandbox / Pilotversuch

Kontrollierte Testumgebung mit begrenztem Nutzer- und Funktionsumfang, in der eine KI-Anwendung vor dem institutionellen Roll-out hinsichtlich Technik, Recht und Nutzerakzeptanz evaluiert wird.

34. Verteilte Kognition

Theorie, dass kognitive Prozesse nicht nur in Individuen, sondern in Interaktion mit Werkzeugen, Artefakten und anderen Personen stattfinden. KI-Systeme fungieren hier als „kognitive Erweiterung“, die Denk- und Problemlöseprozesse maßgeblich beeinflussen kann.

35. Watchlist

Provisorische Liste von KI-Tools, die eine Institution prüfen möchte. Bevor eine offizielle Zulassung oder Einführung erfolgt, werden diese Tools anhand definierter Kriterien (Datenschutz, didaktischer Nutzen, Fairness) bewertet.

9.2 CHECKLISTEN UND VORLAGEN

In diesem Abschnitt werden praxisorientierte Hilfsmittel aufgeführt, die eine konsistente und transparente Umsetzung der KI-Leitlinien unterstützen. Je nach Hochschulkultur und Projektbedarf können diese Vorlagen angepasst oder erweitert werden.

1. Musterformular „KI-Logbuch“

- *Zweck:* Transparente Auflistung sämtlicher KI-Nutzungen (Tools, Einsatzbereiche, Umfang der Überarbeitungen), wie in 6.1 Verantwortung und Transparenz beschrieben.
- *Inhalte:* Eingabefelder für Tool-Name, Version, Einsatzzweck, Datum, Art der Überarbeitung, ggf. Verweise auf Metadaten.

Verwendete Tools	Einsatzbereich	Umfang der Überarbeitungen
------------------	----------------	----------------------------

2. Vorlage „Zulässig vs. unzulässig“

Zweck: Schnellcheck für Lehrende und Studierende, um gängige KI-Anwendungsfälle einzuordnen. Diese Übersicht fasst die zu KI-Nutzungsformen tabellarisch zusammen.

Inhalte: Spalten für KI-Funktion, Einstufung (zulässig / unzulässig), kurzer Hinweis oder Kommentar.

KI-Nutzung	Zulässig?	Anmerkung / Kommentar

3. Dos und Don'ts für Lehrende und Studierende

NUTZUNG VON KI IN DER HOCHSCHULE			
FORUM PRIMAR		Dos & Don'ts für Studierende	
✓	Prüfen Sie zuerst die Kurs- und Prüfungsregeln . Fragen Sie bei Unsicherheiten in der ersten Einheit der Lehrveranstaltung nach, welche Hilfsmittel erlaubt sind.	Davon ausgehen, dass „ alles erlaubt “ ist. Vorgaben variieren; Missachtung gilt möglicherweise als Täuschungsversuch.	✗
✓	KI für Brainstorming, Gliederungen, Grammatik- oder Übersetzungshilfen einsetzen. Nutzen Sie sie als Werkzeug, nicht als Ghostwriter.	KI Kernargumente oder ganze Abschnitte schreiben lassen. Die geistige Leistung muss von Ihnen stammen.	✗
✓	Jede KI-Nutzung offenlegen . Fußnote, Endnote und/oder Eintrag im KI-Logbuch mit Tool, Prompt, Datum und Überarbeitungen.	KI-Einsatz verbergen . Nicht deklarierte Hilfe wird als Täuschungsversuch (Plagiat) behandelt und kann ernste studienrechtliche Konsequenzen nach sich ziehen.	✗
✓	Alles, was die KI ausgibt, sorgfältig verifizieren. Quellen, Zahlen und Zitate in vertrauenswürdiger Literatur überprüfen .	KI-Output ungeprüft übernehmen . LLMs können falsche "Fakten" und erfundene Referenzen erzeugen.	✗
✓	Textbausteine von generativer KI eigenständig umformulieren und redigieren . Bringen Sie Stil, Struktur und Argumentation in Ihre Arbeit.	Roh-Textbausteine der KI unverändert kopieren . Plötzliche Stilbrüche sind leicht erkennbar.	✗
✓	Personen-, Forschungs- und Unternehmens- bzw. Organisations daten schützen . Keine sensiblen Inhalte, wie personenbezogene Daten, in externe Dienste einfügen.	Sensible, personenbezogene und/oder urheberrechtlich geschützte Daten hochladen . Viele Tools speichern Eingaben dauerhaft und/oder verarbeiten sie weiter.	✗
✓	Eingaben respektvoll und professionell formulieren. Beachten Sie Empfehlungen und Richtlinien für diversitätssensibles Verhalten.	Vorurteilsbehaftete, beleidigende oder diskriminierende Inhalte erzeugen. Sie tragen die Verantwortung!	✗
✓	Lernen Sie aktiv die Stärken und Schwächen der eingesetzten KI-Tools kennen , um diese gezielt und kritisch-reflektiert zu nutzen.	Blind auf KI verlassen , ohne deren Funktionsweise und Grenzen verstanden zu haben. Fehlinterpretationen können die Qualität Ihrer Arbeit beeinträchtigen.	✗
✓	Bei Unsicherheit frühzeitig Rat einholen . Der KI-Leitfaden Ihrer Hochschule bietet Ihnen Informationen und den nötigen Rahmen.	Erst nach Abgabe einer schriftlichen Arbeit/eines Arbeitsauftrags offene Fragen klären . Spätere Erklärungen heben Regelverstöße meistens nicht auf.	✗
Die Leitlinien zur Nutzung von KI in der Hochschule geben Ihnen einen Überblick über zentrale Prinzipien für den verantwortungsvollen, transparenten und datenschutzkonforme Einsatz generativer KI. Mit dem QR-Code können Sie einen Chatbot öffnen, der Ihnen weitere Fragen beantwortet. 			

NUTZUNG VON KI IN DER HOCHSCHULE

FORUM PRIMAR

Dos & Don'ts für Lehrende



Sich mit den KI-Regeln der Hochschule vertraut machen.
Die KI-Leitlinien der jeweiligen Hochschule legen fest, welche KI-Nutzung erlaubt ist.



KI vorbildlich und transparent einsetzen.
Nutzen Sie KI bewusst, sichtbar und reflektiert und zeigen Sie ethische Verantwortung.



Offenlegung einfordern.
Fordern Sie von Studierenden dokumentierte Angaben zu verwendeten KI-Beiträgen.



Alle KI-Ergebnisse kritisch überprüfen.
Prüfen und verifizieren Sie KI-generierte Inhalte sorgfältig vor Nutzung und Bewertung.



Datenschutz stets gewissenhaft beachten.
Vermeiden Sie strikt die Eingabe sensibler oder personenbezogener Daten in KI-Tools.



KI-Grenzen vermitteln.
Thematisieren Sie aktiv Chancen, Risiken und mögliche Verzerrungen der KI-Systeme.



Prüfungen KI-resistent und kompetenzorientiert gestalten.
Gestalten Sie Prüfungen mit Transfer-, Reflexions- oder mündlichen Prüfungselementen.



Frühzeitig fachliche KI-Beratung nutzen.
Ziehen Sie rechtzeitig institutionelle Beratung zu didaktischen und rechtlichen Fragen hinzu.



Eigene KI-Kompetenzen regelmäßig erweitern.
Nehmen Sie kontinuierlich an geeigneten Fortbildungen zu KI-Themen teil.

Denken, alles (oder nichts) sei erlaubt.
Vorgaben gelten auch für Lehrende verbindlich.



Aufgaben vollständig an KI delegieren.
Verzichten Sie darauf, Lehre, Bewertung oder Feedback vollständig durch KI zu ersetzen.



KI heimlich oder intransparent einsetzen.
Vermeiden Sie jegliche Form einer nicht offen kommunizierten oder verborgenen KI-Nutzung.



KI-generierte Inhalte ungeprüft verwenden.
Übernehmen Sie keine KI-Inhalte, ohne diese zuvor eigenständig überprüft zu haben.



Weitergabe von geschützten oder sensiblen Daten.
Laden Sie nie vertrauliche, personenbezogene oder urheberrechtlich geschützte Daten hoch.



Blind auf KI vertrauen oder verlassen.
Nutzen Sie KI nur, wenn Ihnen deren Grenzen, Funktionsweise und Verzerrungen bewusst sind.



Prüfungen stellen, die KI einfach lösen kann.
Erstellen Sie keine trivialen Aufgaben, die ohne eigenständiges Denken lösbar wären.



KI-Inhalte ohne Prüfung publizieren.
Veröffentlichen Sie kein KI-generiertes Material ohne vorherige Lizenz- und Rechteprüfung.



KI-Fragen erst nach Abgabe klären lassen.
Klären Sie KI-bezogene Unsicherheiten frühzeitig und niemals erst nach der Abgabefrist.



Die Leitlinien zur Nutzung von KI in der Hochschule geben Ihnen einen Überblick über zentrale Prinzipien für den verantwortungsvollen, transparenten und datenschutzkonformen Einsatz generativer KI. Mit dem QR-Code können Sie einen Chatbot öffnen, der Ihnen weitere Fragen beantwortet.



4. CHATBOT

Ein begleitender Chatbot ermöglicht es, die Empfehlungen für Leitlinien direkt zu befragen und relevante Inhalte bedarfsgerecht aufzurufen.

https://ai.ph-burgenland.at/chatbot/XY7NHfD-0HGQxDN8yJ_us38Kz236Q_JX9yHdswRv8H8



LITERATURVERZEICHNIS:

B

Bundesgesetzblatt für die Republik Österreich. (2024a). *Bundesgesetz, mit dem das Komm-Austria-Gesetz und das Telekommunikationsgesetz 2021 geändert werden* (BGBl. I Nr. 6/2024). <https://www.ris.bka.gv.at/eli/bgbl/I/2024/6>

Bundesgesetzblatt für die Republik Österreich. (2024b). *Bundesgesetz, mit dem das Hochschul-Qualitätssicherungsgesetz geändert wird* (BGBl. I Nr. 25/2024). <https://www.ris.bka.gv.at/eli/bgbl/I/2024/25>

C

Cacho, R. (2024, 1. September). Integrating generative AI in university teaching and learning: A model for balanced guidelines. *Online Learning*, 28(3). <https://doi.org/10.24059/olj.v28i3.4508>

Chan, C. K. Y. (2023, 7. Juli). A comprehensive AI policy education framework for university teaching and learning. *International Journal of Educational Technology in Higher Education*, 20(1). <https://doi.org/10.1186/s41239-023-00408-3>

D

Dabis, A., & Csáki, C. (2024, 6. August). AI and ethics: Investigating the first policy responses of higher education institutions to the challenge of generative AI. *Humanities and Social Sciences Communications*, 11(1). <https://doi.org/10.1057/s41599-024-03526-z>

Datenschutz-Grundverordnung. (2016). *Verordnung (EU) 2016/679 des Europäischen Parlaments und des Rates vom 27. April 2016 zum Schutz natürlicher Personen bei der Verarbeitung personenbezogener Daten, zum freien Datenverkehr und zur Aufhebung der Richtlinie 95/46/EG (Datenschutz-Grundverordnung)*. Amtsblatt der Europäischen Union, L 119, 1–88. <https://eur-lex.europa.eu/legal-content/DE/TXT/?uri=CELEX%3A32016R0679>

Dotan, R., Parker, L. S., & Radzilowicz, J. (2024). Responsible adoption of generative AI in higher education: Developing a “points to consider” approach based on faculty perspectives. In *The 2024 ACM Conference on Fairness, Accountability, and Transparency* (S. 2033–2046). ACM. <https://doi.org/10.1145/3630106.3659023>

E

European Association of Distance Teaching Universities. (2024). *Leading the Future of Learning. Proceedings of the Innovating Higher Education Conference 2024*. <https://doi.org/10.5281/zenodo.14220974>

European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 on artificial intelligence (Artificial Intelligence Act). *Official Journal of the European Union*, L–, 1–144. <http://data.europa.eu/eli/reg/2024/1689/oj>

European Commission. (2025a). *Guidelines on the definition of an artificial intelligence system established by Regulation (EU) 2024/1689 (AI Act)*. Shaping Europe’s Digital Future. <https://digital->

strategy.ec.europa.eu/en/library/commission-publishes-guidelines-ai-system-definition-facilitate-first-ai-acts-rules-application

European Commission. (2025b). *AI Act — Regulatory framework for AI. Application timeline*. Shaping Europe's Digital Future. <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>

H

Hochschulforum Digitalisierung. (2025). *EU AI Act: Wie wird Deutschland KI-kompetent? Herausforderungen und Chancen für die Hochschullehre*. Abgerufen am 8. Juli 2025 von <https://hochschulforumdigitalisierung.de/eu-ai-act-wie-wird-deutschland-ki-kompetent>

Hurenko, O., & Medvedenko, V. (2023). Using artificial intelligence in the educational process: A norm of today or a challenge to academic integrity. *Наукoвi записки БДПУ*, (35), 35–56. <https://doi.org/10.31494/2412-9208-2023-1-3-35-56>

I

ISO. (2023). *ISO/IEC 42001:2023—Information Technology—Artificial intelligence – management system*. International Organization for Standardization. <https://www.iso.org/obp/ui/en/#iso:std:81230:en>

L

Leitgeb, T., Maitz, K., Sitter, G., Matischek-Jauk, M., Möblacher, C., Knaus, M., Gabriel, H., & Meller, S. (2024). *KI-Leitlinien für den PH-Verbund Süd-Ost – Leitlinien für die Nutzung von Künstlicher Intelligenz in der Hochschule*. University College of Teacher Education Burgenland, Eisenstadt. <https://doi.org/10.5281/zenodo.13693466>

Leitgeb, T., & Leitgeb, M. (2024). Empfehlungen für eine evidenzbasierte Nutzung von Künstlicher Intelligenz und Large Language Models in Hochschulen: Ergebnisse eines systematischen Reviews (10.2024). In phpublico: Bd. Einzelbeiträge (S. 1–34). University College of Teacher Education Burgenland. <https://doi.org/10.5281/zenodo.13930620>

Leitgeb, T., Maitz, K., Herbert, G., Möblacher, C., & Matischek-Jauk, M. (2025a). Empfehlungen für KI-Leitlinien des PH-Verbunds Süd-Ost. <https://doi.org/10.5281/zenodo.17426637>

Leitgeb, T., Maitz, K., Matischek-Jauk, M., & Möblacher, C. (2025b). Künstliche Intelligenz im PH-Verbund Süd-Ost: Konzept für eine Untersuchung zum IST-Stand und zu Entwicklungsperspektiven für die Hochschullehre. In L. Hollerer, A. Holzinger, G. Khan-Svik & M. A. Hainz (Hrsg.), *Forschung und Entwicklung in der Primarstufe: Zehn Jahre Forum Primar* (S. 187–194). Leykam Universitätsverlag.

Leitgeb, T., & Leitgeb, M. (2025). Artificial Intelligence and Large Language Models in Higher Education: Results of a Systematic Review. *Ubiquity Proceedings*, 6(1), 33. <https://doi.org/10.5334/uproc.201>

M

Mangal, M., & Pardos, Z. A. (2024). Implementing equitable and intersectionality-aware ML in education: A practical guide. *British Journal of Educational Technology*, 55(5), 2003–2038. <https://doi.org/10.1111/bjet.13484>

Marcel, F., & Kang, P. (2024). Examining AI guidelines in Canadian universities: Implications on academic integrity in academic writing. *Discourse and Writing/Rédactologie*, 34, 93–126. <https://doi.org/10.31468/dwr.1051>

Memarian, B., & Doleck, T. (2024). Human-in-the-loop in artificial intelligence in education: A review and entity-relationship (ER) analysis. *Computers in Human Behavior: Artificial Humans*, 2(1), 100053. <https://doi.org/10.1016/j.chbah.2024.100053>

Mirza, I., Jafari, A. A., Ozcinar, C., & Anbarjafari, G. (2025). *Gender Bias Analysis for Different Large Language Models*. <https://doi.org/10.20944/preprints202501.0016.v1>

Moorhouse, B. L., Yeo, M. A., & Wan, Y. (2023). Generative AI tools and assessment: Guidelines of the world's top-ranking universities. *Computers and Education Open*, 5, 100151. <https://doi.org/10.1016/j.caeo.2023.100151>

Mondal, H., Mondal, S., & Behera, J. K. (2024). Artificial intelligence in academic writing: Insights from journal publishers' guidelines. *Perspectives in Clinical Research*, 16(1), 56–57. https://doi.org/10.4103/picr.picr_67_24

Moya, B., Eaton, S., Pethrick, H., Hayden, A., Brennan, R., Wiens, J., & McDermott, B. (2024). *Academic integrity and artificial intelligence in higher education (HE) contexts: A rapid scoping review*. *Canadian Perspectives on Academic Integrity*, 7(3). <https://doi.org/10.55016/ojs/cpai.v7i3.78123>

Muwagunzi, E., Kabuye, R., Ddamulira, C., & Kizza, S. (2024). Artificial intelligence in academic research at Bugema University: Transforming methodologies and ethical considerations. *East African Journal of Interdisciplinary Studies*, 7(1), 489–503. <https://doi.org/10.37284/eajis.7.1.2484>

O

OECD (2021), *AI and the Future of Skills, Volume 1: Capabilities and Assessments*, Educational Research and Innovation, OECD Publishing, Paris, <https://doi.org/10.1787/5ee71f34-en>.

OECD (2025). Empowering learners for the age of AI: An AI literacy framework for primary and secondary education (Review draft). OECD. Paris. <https://ailiteracyframework.org>

Oh, S. H., & Sanfilippo, M. (2024). *University governance for responsible AI*. In *Proceedings of the ALISE Annual Conference*. University of Illinois Main Library. <https://doi.org/10.21900/j.alise.2024.1706>

P

Perkins, M., & Roe, J. (2024). Academic publisher guidelines on AI usage: A ChatGPT supported thematic analysis. *F1000Research*, 12, 1398. <https://doi.org/10.12688/f1000research.142411.2>

Plata, S., De Guzman, M. A., & Quesada, A. (2023). Emerging research and policy themes on academic integrity in the age of Chat GPT and generative AI. *Asian Journal of University Education*, 19(4), 743–758. <https://doi.org/10.24191/ajue.v19i4.24697>

S

Shchedrina, M. (2024, 30. Juni). *Peculiarities of use of generative artificial intelligence in codes of academic integrity in higher education institutions of Singapore*. *International Scientific Journal of Universities and Leadership*, 17, 154–163. <https://doi.org/10.31874/2520-6702-2024-17-154-163>

Smith, S., Tate, M., Freeman, K., Walsh, A., Ballsun-Stanton, B., Hooper, M., & Lane, M. (2024). *A university framework for the responsible use of generative AI in research*. *arXiv*. <https://doi.org/10.48550/ARXIV.2404.19244>

Spivakovsky, O. V., Omelchuk, S. A., Kobets, V. V., Valko, N. V., & Malchykova, D. S. (2023, 30. Oktober). *Institutional policies on artificial intelligence in university learning, teaching and research*. *Information Technologies and Learning Tools*, 97(5), 181–202. <https://doi.org/10.33407/itlt.v97i5.5395>

T

TÜV Akademie. (2025, 14. April). *Artikel 4 EU AI Act – KI-Kompetenz ist jetzt Pflicht*. Abgerufen am 8. Juli 2025 von <https://die-tuev-akademie.de/blog/ki-verordnung-eu-ai-act-artikel-4-ki-kompetenz-ist-jetzt-pflicht>

U

UNESCO. (2021). *AI in education: Guidance for policy-makers*. UNESCO Publishing. <https://unesdoc.unesco.org/ark:/48223/pf0000376709>

UNESCO. (2022). *Recommendation on the ethics of artificial intelligence*. UNESCO Publishing. <https://unesdoc.unesco.org/ark:/48223/pf0000381137.locale=en>

Universität Graz. (o. D.). *Lehren und Lernen mit KI*. <https://lehren-und-lernen-mit-ki.uni-graz.at/de/>

Urheberrechtsgesetz (UrhG), BGBl. Nr. 111/1936 idF BGBl. I Nr. 244/2021. (2021). *Bundesgesetzblatt der Republik Österreich, Teil I*. <https://www.ris.bka.gv.at/eli/bgbli/2021/244>

W

World Economic Forum. (2024, June 19). *Laut neuem Bericht verlangsamt sich die Dynamik der Energiewende inmitten steigender globaler Volatilität* [Pressemitteilung]. https://www3.weforum.org/docs/WEF_ETI2024_Press_Release_DE.pdf