

FAQ

Künstliche Intelligenz & gute wissenschaftliche Praxis

Version 2

Diese FAQ versammeln Fragen, die uns häufig im Zusammenhang mit künstlicher Intelligenz (KI) und guter wissenschaftlicher Praxis (GWP) erreichen. Die vorliegende Version 2 enthält Updates und Erweiterungen der ursprünglichen Fassung vom November 2024. Die FAQ beschränken sich dabei auf generative KI, insbesondere Large Language Models (LLM), die Forschenden seit der Markteinführung von ChatGPT im November 2022 in verschiedensten Tools zur Verfügung stehen. Die Antworten sollen bei der Orientierung in einem schnelllebigen Thema helfen, ohne dabei präskriptiv zu sein. Sie stellen keine offizielle Positionierung des *Ombudsgremiums für die wissenschaftliche Integrität in Deutschland (OWID)* dar, sondern beschreiben den Status Quo und ordnen bereits bestehende Empfehlungen aus Sicht der GWP ein, identifizieren Lücken und verweisen auf weiterführende Literatur. Diese FAQ richten sich primär an Forschende sowie an Ombudspersonen, die Forschende in Bezug auf Fragen zur KI beraten. Für die Nutzung von KI in der Lehre und in studentischen (Qualifikations-)Arbeiten sind i.d.R. universitäre KI-Richtlinien, angepasste Prüfungsordnungen und Selbstständigkeitserklärungen sowie Entscheidungen individueller Lehrpersonen maßgeblich. Daher werden eventuelle Besonderheiten von KI in der Lehre und in Prüfungsangelegenheiten in diesen FAQ nicht besprochen.

Die Inhalte der FAQ können gern mit entsprechenden Verweis nachgenutzt werden.

Zitierhinweis: Frisch, Katrin (2025). FAQ Künstliche Intelligenz und gute wissenschaftliche Praxis. Version 2. Zenodo.
<https://doi.org/10.5281/zenodo.17349995>

INHALT

1. Zu welchen Aspekten hinsichtlich KI und GWP besteht Konsens?	2
2. Warum sollte die Nutzung von KI angegeben werden?	3
3. Was genau bedeutet „die Nutzung von KI transparent und angemessen angeben“?	3
4. Welche Anwendungen fallen unter den Begriff KI?	3
5. Wo im Manuskript muss die Nutzung von KI angegeben werden?	4
6. Gibt es Zitiervorschläge für die Angabe von KI?	4
7. Müssen Prompts dokumentiert werden?	4
8. Muss die Nutzung von KI für Software deklariert werden?	5
9. Muss die Nutzung von KI angegeben werden, wenn diese nur verwendet wird, um einen Text stilistisch/sprachlich zu verbessern?	5
10. Muss die Nutzung von KI für Übersetzungen angegeben werden?	5
11. Ist die Nutzung von KI zur Inspiration deklarationspflichtig?	6
12. Darf KI für das Schreiben von Anträgen genutzt werden?	6
13. Darf KI zum Generieren von Bildern/Graphiken genutzt werden?	6
14. Darf KI in der Peer-Review eingesetzt werden?	7
15. Was ist bei der Nutzung von KI in Autorschaftsteams zu beachten?	7
16. Was ist der Zusammenhang von KI-generierten Texten und Plagiaten? (Kann ich durch die Nutzung von KI aus Versehen andere Texte plagiiieren?)	7
17. Auf welche Schwächen und Risiken von KI sollten Forschende achten?	8
18. Kann die Nutzung von KI in einem Text durch Software-Tools nachgewiesen werden?	8
19. Ich bin Reviewer:in oder Editor:in und vermute, ein Text bzw. Teile davon wurden mithilfe einer KI geschrieben, diese Nutzung ist jedoch nicht transparent angegeben. Was sollte ich tun?	9
20. Was passiert, wenn ich beschuldigt werde KI genutzt zu haben, ohne dies in der Publikation angegeben zu haben?	9
21. Gibt es aus GWP-Sicht besonders empfehlenswerte KI-Anwendungen?	9
22. Wie ist der Zusammenhang zwischen Urheberrecht und KI-generierten Inhalten?	10
23. Kann ich ausschließen, dass meine Texte als Trainingsmaterial genutzt werden?	10
24. Wie kann ich Konflikten, die aus KI-Nutzung resultieren könnten, vorbeugen?	10
25. Welche KI-Richtlinien sind für mich maßgeblich?	10
26. Warum gibt es keine klaren Empfehlungen vom <i>Ombudsgremium für die wissenschaftliche Integrität</i> zum Thema KI und GWP?	11
Referenzen	12
<i>Literatur</i>	12
<i>Policies, Handreichungen, Empfehlungen</i>	13

1. Zu welchen Aspekten hinsichtlich KI und GWP besteht Konsens?

In Richtlinien und Empfehlungen zu KI hat sich seit der Veröffentlichung von ChatGPT und anderer generativer KI für ein größeres Publikum zu zwei Aspekten ein Konsens herausgebildet.

- 1) KI kann nicht als Autorin fungieren. Dies wird damit begründet, dass KI keine Verantwortung für die Inhalte eines Manuskriptes übernehmen und auch nicht der finalen Fassung zustimmen kann. Beides wird von menschlichen Autor:innen gemäß GWP grundsätzlich gefordert.
- 2) Die Nutzung von KI muss transparent und angemessen im Manuskript angegeben werden.

Die genaue Umsetzung der transparenten Angabe von KI wird jedoch teils unterschiedlich ausgelegt bzw. oft nicht genauer definiert.

In der folgenden Kurzübersicht findet sich eine Zusammenfassung der Editorial Policies in Bezug auf KI der großen Verlage und verwandter Organisationen (Stand: Oktober 2025)

	KI und Autorschaft	Angabe der Nutzung von KI	KI-generierte Abbildungen	KI und Peer-Review	Link zur KI-Policy
Elsevier	KI ≠ Autorin	Nutzung muss im Manuskript angegeben werden	nicht erlaubt	nicht erlaubt	https://www.elsevier.com/de-de/about/policies-and-standards/generative-ai-policies-for-journals
ICMJE	KI ≠ Autorin	Nutzung muss im Manuskript angegeben werden	keine Angabe	sehr eingeschränkt erlaubt	https://www.icmje.org/icmje-recommendations.pdf
Science	KI ≠ Autorin	Nutzung muss im Manuskript angegeben werden	sehr eingeschränkt erlaubt	nicht erlaubt	https://www.science.org/content/page/science-journals-editorial-policies#image-and-text-integrity
Springer Nature	KI ≠ Autorin	Nutzung muss im Manuskript angegeben werden	sehr eingeschränkt erlaubt	sehr eingeschränkt erlaubt	https://www.springernature.com/gp/policies/editorial-policies
Taylor & Francis	KI ≠ Autorin	Nutzung muss im Manuskript angegeben werden	nicht erlaubt	sehr eingeschränkt erlaubt	https://taylorandfrancis.com/our-policies/ai-policy
Wiley	KI ≠ Autorin	Nutzung muss im Manuskript angegeben werden	keine Angabe (Journal Policy) eingeschränkt (Buch Policy)	sehr eingeschränkt erlaubt	Journal Policy https://authorservices.wiley.com/ethics-guidelines/index.html#22 Buch Policy https://www.wiley.com/en-us/publish/book/resources/ai-guidelines/
WAME	KI ≠ Autorin	Nutzung muss im Manuskript angegeben werden	keine Angabe	sehr eingeschränkt erlaubt	https://wame.org/page2.php?id=106

2. Warum sollte die Nutzung von KI angegeben werden?

Die Angabe der Nutzung von KI entspricht den geltenden Transparenzstandards der GWP (siehe [Leitlinien 12](#) und [13](#) im DFG-Kodex „Leitlinien zur Sicherung guter wissenschaftlicher Praxis“). Angaben zur Nutzung von KI erlauben es Lesenden und Reviewer:innen, die Ergebnisse, Methoden und Arbeitsschritte besser nachzuvollziehen und einzuordnen. Aufgrund der Vielzahl von unterschiedlichen KI-Anwendungen und Funktionen gibt es bzgl. der genauen Gestaltung der Angaben jedoch unterschiedliche Empfehlungen.

3. Was genau bedeutet „die Nutzung von KI transparent und angemessen angeben“?

Darauf gibt es keine einheitliche Antwort. Die meisten Editorial Policies spezifizieren die Angabe der Nutzung nicht oder nur minimal. Vorschläge zur Angabe von KI in der Literatur unterscheiden sich in den Details bzw. der Komplexität der Angabe, allen gemein ist jedoch ein gewisser Minimalkonsens. Dazu gehört die Angabe

- der genutzten KI-Anwendung, inklusive Version, Datum der Nutzung, URL und
- wofür die KI-Anwendung wie genutzt wurde.

Einige Policies bieten jedoch detaillierte Informationen zur Deklaration. Die [AI Policy von Wiley](#), die sich speziell an Buchautor:innen richtet, listet beispielsweise, welche Nutzungsarten deklariert werden müssen und welche Angaben Autor:innen machen müssen. Zudem werden darin Beispieldeklarationen angegeben. Es gibt auch Vorstöße von kleineren Verlagen, z.B. [Berlin Universities Publishing \(BUP\)](#), die [konkrete Vorschläge erarbeitet haben](#). Einige Verlage erwarten mittlerweile außerdem eine Bewertung oder Reflexion der KI-Nutzung bzw. KI-generierten Inhalte durch die Autor:innen (z.B. [Wiley](#), [PLOS One](#)). Forschende, die sich mit der Thematik befassen, schlagen zudem vor, anzugeben, welche Person aus einem Team für die Nutzung der KI verantwortlich war (siehe [Hosseini et al. 2023](#)) oder ausführliche Überlegungen zur technischen Arbeitsweise der genutzten KI und ihren Limitierungen in die Deklaration einfließen zu lassen ([Resnik und Hosseini 2024](#)). Insbesondere letzterer Vorschlag geht weit über grundsätzlich geforderte Angaben hinaus, kann aber durchaus als Denkanstoß im Umgang mit KI dienlich sein, um zu prüfen, ob das favorisierte KI-Tool zur beabsichtigten Nutzung passt, und auf welche Schwächen geachtet werden muss. Beides setzt eine intensive Auseinandersetzung mit dem verwendeten Tool voraus.

Andere Ansätze zur Deklaration bieten das Artificial Intelligence Disclosure (AID) Framework sowie das AI Attribution Toolkit. Das [Artificial Intelligence Disclosure \(AID\) Framework von Kari D. Weaver](#) orientiert sich an der bestehenden Credit Roles Taxonomy (CReDiT). Dabei wird in einem kurzen Statement ausgewiesen, welche Rolle(n) das genutzte KI-Tool im Forschungsprozess übernommen hat. Das [AI Attribution Toolkit](#) – entwickelt von Mitarbeitenden von IBM Research ausgehend von [He et al. 2025](#) – wurde inspiriert von den Creative Commons Lizenzen und fokussiert sich insbesondere darauf, das proportionale Verhältnis von menschlichen versus KI-Beiträgen abzubilden. Es soll Forschenden ermöglichen in konziser Weise nicht nur die Form der KI-Nutzung anzugeben, sondern auch zu spezifizieren, wie viel generative KI zum Einsatz kam.

Generell sollte die Angabe der Nutzung den Transparenzbedürfnissen der Leser:innen und Reviewer:innen gerecht werden, aber auch der tatsächlichen Arbeitsweise mit KI entsprechen. Beides kann sehr fachspezifisch ausfallen und sollte deswegen auch mit der eigenen Community besprochen werden. Dies betrifft auch die Frage, welche Nutzungsarten und KI-Anwendungen der Dokumentationspflicht unterliegen ([siehe dazu Fragen 8-11](#)).

4. Welche Anwendungen fallen unter den Begriff KI?

Der Begriff Künstliche Intelligenz ist nicht eindeutig definiert und die Nutzung des Begriffes für verschiedene Anwendungen auch nicht gänzlich unumstritten. Im allgemeinen Sprachgebrauch dient in

der neueren Debatte um künstliche Intelligenz meist der Chatbot „ChatGPT“ als Synonym für generative künstliche Intelligenz, insbesondere für Large Language Models (LLMs). Generative KI-Anwendungen, die für Forschende relevant sind, gibt es derweil viele (siehe z.B. [Liste von KI-Ressourcen zusammengestellt durch VK:KIWA](#) oder den [Ithaka Product Tracker](#)). Im Zuge der steten technischen Weiterentwicklung lässt sich eine zunehmende Integration von KI-Technologien in bestehende Anwendungen beobachten. Eine klare Grenzziehung zwischen (deklarierungspflichtigen) KI-Anwendungen und anderen (nicht deklarierungspflichtigen) Softwaretools ist deswegen teils schwierig. In Richtlinien für Forschende wird sich momentan zumeist auf generative KI, insbesondere LLMs, beschränkt. KI-Anwendungen, die nur zur Sprachbearbeitung/-verbesserung genutzt werden, werden meist nur als Hilfsmittel klassifiziert ([siehe Frage 9](#)). Spezialisierte KI-Anwendungen, die bei den ersten Schritten eines Forschungsprojektes unterstützen (z.B. Hypothesenfindung oder Literatur-Review), werden nicht immer in Empfehlungen dezidiert mitgenannt. Eine Ausnahme bilden die [Living Guidelines on the Responsible Use of Generative AI in Research](#) der Europäischen Kommission, in denen diese Nutzungsarten als potentiell „substantial use“ klassifiziert werden und demnach die verwendeten KI-Tools deklarationspflichtig wären. Einige Verlags-Policies grenzen deutlich ab, auf welche KI sich welche Regelungen und Transparenzangaben beziehen. Die [Policy von Elsevier](#) beispielsweise unterscheidet zwischen KI-Anwendungen für den Schreibprozess, für den Forschungsprozess sowie weiteren assistierenden Anwendungen: Für alle drei gibt es unterschiedliche Regelungen. Autor:innen sollten sich im Vorfeld mit den für sie maßgeblichen Regelungen vertraut machen und im Zweifelsfall eher den größtmöglichen Transparenzbedingungen gerecht werden.

5. Wo im Manuskript muss die Nutzung von KI angegeben werden?

Dazu gibt es in den Richtlinien und Empfehlung nicht immer genaue Angaben. Oft wird jedoch der Methodenteil oder die Danksagung als Ort für die Angabe der KI-Nutzung vorgeschlagen. Auch gibt es Vorschläge, die Deklaration als separates Statement am Ende des Manuskripts anzuführen. Welcher Teil des Manuskriptes sich für die Angabe anbietet, kann auch von fachspezifischen Konventionen sowie der Form der Angabe abhängen. So kann zum Beispiel eine sehr detaillierte Angabe der KI-Nutzung inklusive Prompts und Chatverläufen auch als Supplement hinterlegt werden (wie z.B. [ASC Nano](#) vorschlägt).

6. Gibt es Zitiervorschläge für die Angabe von KI?

Sofern institutsintern oder vom Journal keine konkreten Zitationsvorschläge zur Angabe von KI vorgegeben sind, können bestehende Vorschläge etablierter Zitationsstile zur Orientierung genutzt werden. Dazu gehören die [American Psychological Association](#) (APA), das [Chicago Style Manual](#) sowie die [Modern Language Association](#) (MLA). Dabei sollte darauf geachtet werden, dass die Angabe alle wichtigen Elemente enthält, die vom Zielmedium gefordert werden.

7. Müssen Prompts dokumentiert werden?

Dazu gibt es keine einheitlichen Angaben. Von den großen Verlagen macht nur [Science](#) genaue Vorgaben, demnach sollen Prompts im Methodenteil angegeben werden. Die [APA](#), das [Chicago Style Manual](#) sowie die [MLA](#) geben an, dass Prompts im Manuskript angeführt werden können, schreiben dies aber nicht vor. Die APA empfiehlt Forschenden ihre genutzten Prompts für die interne Dokumentation aufzuheben. Unter Forschenden wird die Nützlichkeit der Angabe von Prompts z.T. hinterfragt. Dies wird damit begründet, dass eine solche Angabe nicht der tatsächlichen (oft iterativen) Arbeitsweise mit KI entsprechen würde. Hinzu kommt, dass selbst bei gleichlautenden Prompts aufgrund der stochastischen Funktionsweise keine deckungsgleichen Antworten erfolgen.

Die Frage, ob Prompts angegeben werden müssen, sollte daher weiterhin insbesondere von den Fachcommunitys diskutiert werden. Zu überlegen wäre, welche Funktion die Angabe der Prompts für die Fachcommunity bzw. für die jeweilige Publikation erfüllt. Auch wenn Ergebnisse der KI nicht

reproduzierbar sind, können die genauen Prompts, die eingesetzt wurden, ggf. Aufschlüsse über die Arbeitsweise der Autor:innen geben. Diese Form der Transparenz kann aber auch über andere Angaben erfolgen, z.B. eine Reflexion über den Einsatz von KI für das betreffende Manuskript (siehe [Frage 3](#)).

8. Muss die Nutzung von KI für Software deklariert werden?

In den Editorial Policies der Verlage wird meist nicht dezidiert auf die Nutzung von generativer KI für Software eingegangen. Momentan macht nur die [AI Policy für Buchautor:innen von Wiley](#) dazu Angaben und spezifiziert, dass das Generieren und Modifizieren von Software-Code mit Hilfe von KI deklarationspflichtig ist. Die [World Association of Medical Editors empfiehlt](#), dass KI-Tools, die bei der Erstellung von Code beteiligt waren, im Manuskript angegeben werden sollten. Aus Transparenzgründen sollte generell für Software gelten, dass die Nutzung von KI bei der Erstellung und Überarbeitung des Codes transparent angegeben werden sollte. Dies ist insbesondere dann wichtig, wenn die Software anderen Forschenden zur Verfügung gestellt bzw. publiziert wird.

9. Muss die Nutzung von KI angegeben werden, wenn diese nur verwendet wird, um einen Text stilistisch/sprachlich zu verbessern?

Einige Editorial Policies unterscheiden zwischen verschiedenen Verwendungsarten von KI bzw. zwischen verschiedenen KI: meist in generative KI auf der einen und Anwendungen zum Überprüfen von Grammatik und Stil, wie z.B. Grammarly, auf der anderen Seite. Die Richtlinien zur Angabe von KI beziehen sich dann meist nicht auf Letztere (siehe beispielsweise die [AI Policy von Elsevier](#) oder [Wiley](#)). So spezifiziert Wiley in seiner KI-Policy für Buchautor:innen, dass „[b]asic grammar or spell checking“ sowie „[s]imple language polishing“ nicht deklarationspflichtig sind. Auch im Rahmen der [Stellungnahme der DFG](#) werden KI-Anwendungen, „die sich nicht auf den wissenschaftlichen Inhalt des Antrags [auswirken] (bspw. Grammatik-, Stil-, Rechtschreibprüfung, Übersetzungsprogramme)“, von der Kennzeichnungspflicht ausgenommen. Ebenso plädieren Stimmen aus der Forschungscommunity dafür, dass die Angabe der Nutzung von generativer KI als Schreibassistent freiwillig sein sollte ([Hosseini et al. 2025](#)).

Innerhalb der Fächer kann diese Frage auch unterschiedlich bewertet werden – je nachdem, welche Stellung Text bzw. sprachliche Formulierungen innehaben. Insbesondere in den Geisteswissenschaften ist der Sprachstil eng an Personen bzw. Schulen geknüpft. In diesem Fall wäre die Angabe der Nutzung von Anwendungen, die den Sprachstil verbessern, denkbar.

10. Muss die Nutzung von KI für Übersetzungen angegeben werden?

KI-generierte Übersetzungen werden in den Editorial Policies und Empfehlungen oft nicht dezidiert adressiert. Die wenigen Ausnahmen machen jedoch ganz unterschiedliche Vorgaben. Die [KI Policy für Buchautor:innen von Wiley](#) sieht Übersetzungen als deklarationspflichtig an, wohingegen die Empfehlung der DFG Übersetzungsprogramme als Beispiel für nicht deklarationspflichtige KI-Anwendungen anführt (siehe [Frage 9](#)). Auch die [BUP-Handreichung zur Zitation von KI-Tools](#) klassifiziert Übersetzungstools als Hilfsmittel, wonach „nicht zwingend eine explizite Nennung“ erforderlich ist.

Autor:innen sollte bewusst sein, dass beim maschinellen Übersetzen von Texten wichtige Informationen verloren gehen oder verzerrt werden können. Auch hier gilt, dass die übersetzten Texte sorgfältig geprüft werden müssen und Autor:innen die Verantwortung für etwaige Fehler übernehmen. Übersetzungen können außerdem als wichtige Eigenleistung und Skill bzw. Selbstverständnis in einem Fach (z.B. den fremdsprachlichen Philologien) angesehen werden. Insbesondere der Übersetzung von literarischen Texten kommt ein hoher Stellenwert zugute: „Eine Übersetzung ist das Ergebnis eines individuellen Umgangs mit einem Ausgangswerk. Dieser Umgang muss gewissenhaft verantwortet werden, nicht nur im eigenen Namen, sondern auch in dem des Autors bzw. der Autorin des Originals“ ([VdÜ / A*ds / IGÜ – Offener Brief zur KI-Verordnung](#)). In Fächern wie den Philologien und bei der

Übersetzung von literarischen Texten (für die wissenschaftliche Arbeit) wäre die Angabe der Nutzung von Übersetzungstools empfehlenswert.

11. Ist die Nutzung von KI zur Inspiration deklarationspflichtig?

Abhängig davon, wie man „Inspiration“ definiert, gibt es zu dieser Frage teils konkrete Empfehlungen. [BUP empfiehlt](#) bei der Nutzung von KI als rudimentäre Inspirationsquelle einen allgemeinen Hinweis auf die Nutzung als „Anmerkung oder ggf. im Methodenteil“. In den [Living Guidelines on the Responsible Use of Generative AI in Research](#) der Europäischen Kommission findet sich die Formulierung, dass der substantielle Gebrauch von KI-Tools deklariert werden sollte – wozu u.a. „identifying research gaps“ und „developing hypotheses“ gezählt werden können (S. 7-8). Generell gilt es auch hier zu prüfen, ob im anvisierten Journal/Verlag dazu Vorgaben gemacht werden. Fehlen konkrete Vorgaben, sollte eine Entscheidung mit Blick auf fachspezifische Leseerwartungen sowie auf die Bedeutung und den Umfang der KI-generierten Ideen getroffen werden.

12. Darf KI für das Schreiben von Anträgen genutzt werden?

Hier sind die Angaben der jeweiligen Forschungsförderer zu beachten. Die Deutsche Forschungsgemeinschaft (DFG) erklärt in ihrer [Stellungnahme zum Einfluss generativer Modelle für die Text- und Bilderstellung auf die Wissenschaften und das Förderhandeln der DFG](#) (2023), dass die Nutzung von KI in Anträgen zulässig sei, sofern diese klar markiert ist ([siehe auch Frage 3](#)). Ferner führt die Stellungnahme aus, dass die Nutzung von KI in Anträgen neutral bewertet wird, also weder negativ noch positiv seitens der Gutachter:innen geltend gemacht werden soll. Zu beachten ist, dass die Stellungnahme die Nutzung von KI im Rahmen der Begutachtung hingegen verbietet ([siehe auch Frage 14](#)).

13. Darf KI zum Generieren von Bildern/Graphiken genutzt werden?

Wie in der obigen Verlagsübersicht einsehbar, haben Journals dazu häufig sehr restriktive Vorgaben. Meist ist die Nutzung nur in einem sehr begrenzten Rahmen erlaubt, z.B. in thematischen Beiträgen, die sich dezidiert mit dem Thema KI auseinandersetzen. Die Nutzung muss dann ebenfalls ausreichend gekennzeichnet sein. Ausnahmen sind z.B. der Verlag Frontiers, der die Nutzung von KI zum Erstellen von Bildern generell erlaubt, mit dem Verweis, dass die Nutzung klar gekennzeichnet werden muss ([siehe Frontiers Artificial intelligence: fair use and disclosure policy](#)). Die detaillierteste AI Policy zum Thema ist momentan die von [Wiley für Buch-Autor:innen](#): Sie listet auf, welche konkreten Nutzungsarten in Bezug auf KI-generierte Bilder erlaubt und welche nicht gestattet sind. Sie liefert darüber hinaus Hinweise, welche Anforderungen erlaubte KI-generierte Bilder (dazu gehören z.B. Flussdiagramme und Visualisierungen) erfüllen müssen. Das Veröffentlichen von KI-generierten Bildern, die als echte Forschungsdaten, wie beispielsweise Western Blots, Zellproben, Präparate oder Artefakte, ausgegeben werden, ist nicht erlaubt. Aus Sicht der GWP stellen sie einen Akt der Täuschung dar und können als Datenfabrikation (und damit als wissenschaftliches Fehlverhalten) gewertet werden.

Für die Nutzung von KI für die Bildgeneration bei anderen wissenschaftlichen Formaten, z.B. Vorträge oder Poster-Präsentationen, existieren momentan keine übergreifenden Vorgaben. Hier sollte zuerst mit der eigenen Projektgruppe/Fachcommunity gesprochen werden. Auch diesbezüglich gilt, dass KI-generierte Bilder nicht zur Vortäuschung von erhobenen Forschungsdaten bzw. -ergebnissen genutzt werden dürfen. Bei Bildern, die zu rein illustrativen Zwecken auf Vortragsfolien genutzt werden, spricht aus Sicht der GWP nichts gegen eine Nutzung von KI. Bei Graphiken, die Forschungsprozesse illustrieren (z.B. Diagramme oder Schemata), sollte die Nutzung von KI möglich sein, ggf. mit Deklaration. Wie bei KI-generiertem Text obliegt es den Autor:innen, das Ergebnis auf Korrektheit zu überprüfen. Zusätzlich zur generellen Frage der Nutzung von KI zum Generieren von Abbildungen für Präsentationen und Poster ist momentan auch die damit zusammenhängende Deklaration noch nicht geregelt. Es

empfiehlt sich auf bestehende Vorschläge zur Nutzungsangabe zurückzugreifen. Aufgrund der kurzen übersichtlichen Darstellung sei an dieser Stelle besonders auf das [AI Attribution Toolkit](#) verwiesen ([siehe auch Frage 3](#)).

14. Darf KI in der Peer-Review eingesetzt werden?

Für die Begutachtung ist die Nutzung von KI in vielen Richtlinien und Empfehlungen entweder starken Einschränkungen unterworfen oder gar nicht erlaubt ([siehe Verlagsübersicht unter Frage 1](#)). Generell gilt, dass die Einspeisung des zu begutachtenden Manuskriptes (oder Projektantrags) in eine KI nicht erlaubt ist. Angeführt werden dafür die Aspekte Vertraulichkeit und Datenschutz, die dadurch ggf. kompromittiert werden. Auch sollte bedacht werden, dass die Funktion der Begutachtung der Wissensproduktion und Weiterentwicklung des eigenen Faches dient. Diese lenkenden Aufgaben sollten nicht an eine KI ausgelagert werden ([Hosseini und Horbach 2023](#), siehe auch [Bergstrom und Bak-Coleman 2025](#)). Der Einsatz von generativer KI in der Peer Review wird jedoch fortlaufend diskutiert. Laut einer [aktuellen Umfrage von Nature](#) sieht die Mehrheit der befragten Forschenden den Einsatz von generativer KI bei der Einschätzung eines zu begutachtenden Manuskripts als kritisch; eine Assistenzfunktion, z.B. zur Beantwortung von Fragen zum Manuskript, war für über die Hälfte der Befragten jedoch denkbar. Obwohl der Einsatz von KI für Peer Review momentan von Verlagen nicht zugelassen ist, zeigen [Studien](#), dass eine Minderheit der Befragten trotzdem davon Gebrauch macht. Die Verlage selbst nutzen mittlerweile oft verschiedene KI-Tools zum Pre-Screening von eingesandten Manuskripten.

KI darf, wenn überhaupt, dann nur zur sprachlichen Nachbearbeitung des Reviews genutzt werden ([siehe Verlagsübersicht unter Frage 1](#)). Reviewer:innen sollten jeweils prüfen, welche Vorgaben für sie gelten.

15. Was ist bei der Nutzung von KI in Autorschaftsteams zu beachten?

Aufgrund der relativen Neuheit vieler KI-Anwendungen sollte in Autorschaftsteams zu Projektbeginn geklärt werden, ob alle Autor:innen mit der generellen Nutzung von KI-Tools und dem Umfang der Nutzung einverstanden sind. Insbesondere bei trans- und interdisziplinären Autorschaftsteams ist dies wichtig, da hier z.B. unterschiedliche Bewertungen von Textarbeit aufeinandertreffen (siehe auch [Frage 9](#)). In Teams bietet sich außerdem eine gute, interne Dokumentation der Nutzung von KI an. In Fällen von nachträglich auftretenden Konflikten oder GWP-Verstößen kann der Prozess dann besser nachvollzogen werden. Autor:innen, die in Teams publizieren, können zudem überlegen, ob sie der Zitationsempfehlung von [Hosseini et al. 2023](#) folgen, die bei der Angabe der Nutzung von KI zusätzlich die anwendende Person notiert ([siehe auch Frage 3](#)).

16. Was ist der Zusammenhang von KI-generierten Texten und Plagiaten? (Kann ich durch die Nutzung von KI aus Versehen andere Texte plagieren?)

Da generative KI-Modelle auf Grundlage einer riesigen Anzahl an Trainingsdaten stochastisch arbeiten, gehört unbeabsichtigtes Plagieren nicht zu den häufigen Risiken von KIs. Dies stützt sich auf ein Verständnis von Plagiaten im Sinne einer ungekennzeichneten Wiederverwendung von Beiträgen Dritter. Diese Definition setzt voraus, dass es eine zuordenbare Quelle oder Person gibt, aus der sich wortwörtlich oder mit Paraphrase bedient wurde, ohne dass diese genannt wird. Es gibt Beispiele, die aufzeigen, dass bekannte KI-Tools trotz ihrer stochastischen Natur bestimmte Texte fast wortwörtlich reproduzieren können ([Cooper et al. 2025](#), [Henderson et al. 2023](#) oder der Fall New York Times versus OpenAI, siehe z.B. [Pope 2024](#)). In den Tests wurde per Prompt jedoch genau auf dieses Resultat hingearbeitet.

In bestimmten Fächern der Geisteswissenschaften, in denen Text und Sprache eine übergeordnete Rolle einnehmen, können jedoch schon einzelne Begrifflichkeiten bestimmten

bekannten Theoretiker:innen zugeordnet werden. Fehlt im generierten Text die Ausweisung, auf wen der Begriff zurückgeht, kann dies daher entweder als unzureichende Referenzierung oder gar als Plagiat angesehen werden (siehe Seadle „For the humanities, words matter. [...] A stolen word is a stolen thought“ (42)). Forschende, die sich länger in einer Disziplin bewegen, sind sich der typischen Fachdiskurse zumeist bewusst; deswegen sind besonders interdisziplinär Forschende angehalten, KI-generierte Inhalte gründlich zu prüfen. Generell sollte immer eine umfassende Überarbeitung von KI-generiertem Text durch die Forschenden erfolgen.

Eine zweite Einschränkung, die hier kurz erwähnt sei, ist, dass im prüfungsrechtlichen Sinne in Hochschulen auch ein weiteres Verständnis von Plagiaten angewendet wird. So schreibt zum Beispiel Ulrike Verch, dass eine Übernahme von KI-generiertem Output "ohne Kennzeichnung und umfassende Überarbeitung ein Plagiat begründen [könnte], da die Prüfungsordnungen der Hochschulen in der Regel voraussetzen, dass die eingereichten Leistungen eigenständig ohne andere als die angegebenen und erlaubten Hilfsmittel erbracht worden sind" ([Verch 2023, Seite 11](#)).

17. Auf welche Schwächen und Risiken von KI sollten Forschende achten?

Für generative KI ist besonders das sogenannte Halluzinieren eine bekannte Schwäche. Es gibt jedoch eine ganze Reihe von Risiken und Schwächen, die KI-Anwendungen mit sich bringen. Dazu gehören fehlende oder falsche Referenzen, Fehler in wörtlichen Zitaten, komplett fabrizierte Zitate oder Referenzen, eine Verzerrung des wiedergegebenen Diskurses insbesondere bei kontroversen Themen, die Weitergabe veralteter Informationen sowie die Reproduktion von Bias bzw. vorurteilsbehafteten Aussagen. Der große Anteil englischsprachiger (oft auch US-amerikanischer) Quellen im Trainingsmaterial lädt zu Vorsicht bei der Nutzung von KI in anderen Sprachen ein. Eine Typologie der Schwächen und Risiken von generativen KI-Anwendungen findet sich beispielsweise bei [Oertner 2024](#). Viele generative KI-Tools weisen mittlerweile explizit daraufhin, dass generierter Output Fehler enthalten kann. Forschende sollten jedoch auch bei Tools, die vorgeben besonders zuverlässige Ergebnisse zu liefern, vorsichtig sein und jeglichen Output prüfen. Autor:innen übernehmen die Verantwortung für die von KI produzierten Falschangaben und Regelverstöße. Durch gutes Prompting und gutes Prüfen lassen sich die Risiken der Nutzung minimieren. Dazu gehört ein gutes Verständnis über die Arbeitsweise und Funktionen des genutzten KI-Tools.

Auf höherer Ebene betrachtet bergen generative KI-Tools weitere mögliche Risiken. Dazu gehören epistemische Risiken (siehe z.B. [Messerli und Crockett 2024](#), [Schütze 2025](#)), De-Skilling, die Zunahme von Literatur aus Paper Mills, die Überlastung von wissenschaftlichen Infrastrukturen durch Bots (z.B. [Weinberg 2025](#)), ein Verlust von digitaler Autonomie ([Bahr 2024](#)), die Abhängigkeit von proprietären Diensten bzw. kommerziellen Anbietern, Urheberrechtsverletzungen sowie die Ausbeutung von Menschen und natürlichen Ressourcen (Hao 2025). Eine umfassende Liste von Risiken, die im Zusammenhang mit KI auftreten können, findet sich im [MIT AI Risk Repository](#).

18. Kann die Nutzung von KI in einem Text durch Software-Tools nachgewiesen werden?

Es gibt mehrere Studien, die sich mit Detektions-Tools befassen haben, die teilweise zu unterschiedlichen Ergebnissen hinsichtlich des genauen Potentials von Software-Tools kommen ([Weber-Wulff et al. 2023](#), [Gao et al. 2023](#)). Generell lässt sich mit Blick auf die Studien jedoch festhalten, dass es keine hinreichend zuverlässigen Tools zur Erkennung von KI-generierten Inhalten auf dem Markt gibt. Auch menschliche Reviewer:innen konnten in Studien nur unzureichend KI-generierte Texte von menschlich-geschriebenen unterscheiden. Von einer eindeutigen Feststellung der Nutzung von KI in Texten ist im Moment also nicht auszugehen.

19. Ich bin Reviewer:in oder Editor:in und vermute, ein Text bzw. Teile davon wurden mithilfe einer KI geschrieben, diese Nutzung ist jedoch nicht transparent angegeben. Was sollte ich tun?

Das kommt darauf an, was genau diese Vermutung auslöst. Es gibt einige deutliche Anzeichen, dass eine KI zum Einsatz kam, z.B. gängige Phrasen in Texten („Certainly, here is an introduction for you“, „As an AI language model, I cannot...“, „as of my last knowledge update“), Nonsense-Wörter oder verzerrte Schriftzeichen bei der Erstellung von Bildern. Hierbei würde es sich um einen belegten und somit eindeutigen Verstoß gegen die gängigen Vorgaben bzgl. der KI-Nutzung handeln. Wie ein Verstoß gegen die KI-Policy gehandhabt wird, sollte journalintern festgelegt werden. Einige Verlagspolicies spezifizieren rudimentär das Vorgehen, falls nicht-deklariert KI-Inhalt in Manuskripten vermutet wird. Ist die Vermutung auf weniger konkrete Anzeichen zurückzuführen und macht sich eher am Sprachstil generell fest bzw. an der Nutzung von bestimmten Wörtern, die mittlerweile oft als Signalwörter für KI-Nutzung angesehen werden (z.B. „delve“, „meticulous“, „commendable“), sollte ein klärendes Gespräch mit den Autor:innen gesucht werden, denn weder die genannten Signalwörter noch eine Prüfung mit einer Software kann hier zweifelsfrei die Nutzung von KI belegen. Eine vorschnelle Anschuldigung eines Verstoßes gegen die GWP sollte mit Blick auf [Leitlinie 18](#) des DFG-Kodex vermieden werden, in der es heißt: „[b]ewusst unrichtig oder mutwillig erhobene Vorwürfe können selbst ein wissenschaftliches Fehlverhalten begründen“.

20. Was passiert, wenn ich beschuldigt werde KI genutzt zu haben, ohne dies in der Publikation angegeben zu haben?

Dazu gibt es erste Erfahrungsberichte ([Wolkovich 2024](#)). Wie in [Frage 18](#) aufgeführt gibt es keine Software-Tools, die die Nutzung von KI sicher nachweisen können. Autor:innen sollten sich die genaue Begründung geben lassen, wieso ihr Text bzw. Teile davon im Verdacht stehen KI-generiert zu sein. Hilfreich ist es, den Arbeits- und Schreibprozess für sich selbst intern zu dokumentieren, damit anhand der Dokumentation ggf. nachvollzogen bzw. nachgewiesen werden kann, dass nicht mit KI gearbeitet wurde.

21. Gibt es aus GWP-Sicht besonders empfehlenswerte KI-Anwendungen?

Welche KI-Anwendung die geeignetste für eine Aufgabe ist, hängt von verschiedenen Faktoren und deren Gewichtung ab. So sollten Forschende neben der Bewertung nach Funktionen und Leistung von KI-Anwendungen auch prüfen, ob die gewählte KI-Anwendung gesetzliche Vorschriften und Bestimmungen einhält. Datenschutz und Vertraulichkeit spielen dabei eine übergeordnete Rolle. Viele KI-Anwendungen sind außerdem dafür bekannt, dass sie Schwächen und Risiken bzgl. der generierten Inhalte mit sich bringen ([siehe Frage 17](#)). Autor:innen sollten sich dieser bewusst sein und die generierten Inhalte daraufhin sorgfältig prüfen.

Die Entscheidung mit KI oder einer bestimmten KI zu arbeiten, sollte deswegen im Vorfeld mit Bedacht getroffen werden. Bevor ein Tool für das wissenschaftliche Arbeiten eingesetzt wird, sollten sich Forschende ausgiebig damit auseinandersetzen. Dies gilt für ethische Fragen im Zusammenhang mit bestimmten KI-Anwendungen sowie KI-Anwendungen im Allgemeinen (siehe ebenfalls Frage 17). Auch eine Auseinandersetzung mit technischen Aspekten von KI-Tools (Trainingsdaten, Model Weights, Model Cards) kann bei der Bewertung hilfreich sein. Proprietäre Modelle erlauben jedoch nur begrenzt Einblick in technische Aspekte. Hier könnte eine Entscheidung für offene (bzw. offenere) Modelle Abhilfe schaffen. Einen Überblick gibt der [European Open Source AI Index](#). Eine Art Leitfaden zur Auseinandersetzung mit KI aus Sicht der GWP können dabei die Überlegungen von [Resnik und Hosseini 2024](#) bieten. Auch der [MLA Guide to AI Literacy](#) (der sich zwar an Studierende richtet, aber zum Großteil auf Forschung übertragbar ist) kann in dieser Hinsicht als hilfreicher Leitfaden dienen.

22. Wie ist der Zusammenhang zwischen Urheberrecht und KI-generierten Inhalten?

Fragen im Zusammenhang mit Urheberrecht und KI-Anwendungen können sowohl die generierten Inhalte betreffen (z.B. Habe ich das Urheberrecht an den von mir mithilfe einer KI-generierten Inhalten?) als auch die in eine KI eingespeisten Inhalte (Kann ich das Urheberrecht Anderer durch meine Arbeit mit KI verletzen?). Mit diesen Fragen haben sich bereits [Roman Konertz 2023](#) und [Ulrike Verch 2024](#) sowie die [FAQ vom Börsenverein des Deutschen Buchhandels](#) eingehend beschäftigt, auf deren Publikationen wir an dieser Stelle verweisen.

Fragen zum Urheberrecht bzw. Copyright stellen sich auch im Hinblick auf das Training von LLMs. Ob diese Praxis unter die Fair Use Prinzipien des Copyrights fällt, ist derzeit Gegenstand von Gerichtsprozessen in den USA, deren Ergebnisse abgewartet werden müssen. Eine erste Entscheidung wurde im Juni 2025 im Fall Anthropic (Firma des KI-Tools Claude) versus Autor:innen Andrea Bartz, Charles Graeber und Kirk Wallace gefällt. Dabei ging es um das Training an urheberrechtlich geschütztem Material, das unter anderem aus sogenannten Schattenbibliotheken stammte. Das Urteil sah das Training an dem urheberrechtlich geschützten Material als fair use an, nicht jedoch das Herunterladen von Raubkopien zum Erstellen einer Trainingsbibliothek ([Bartz v. Anthropic PBC, 3:24-cv-05417, \(N.D. Cal.\)](#)).

23. Kann ich ausschließen, dass meine Texte als Trainingsmaterial genutzt werden?

Bei urheberrechtlich geschützten Werken bedarf es theoretisch einer Zustimmung der Rechteinhaber:innen, wenn die Werke zum Training von LLMs genutzt werden sollen (siehe [FAQ Börsenverein des Deutschen Buchhandels](#)). Forschende sollten sich informieren, ob Verlage, in denen sie veröffentlichen, möglicherweise bereits Lizenz-Verträge mit KI-Firmen abgeschlossen haben. Bei Open Access Texten, die unter CC-Lizenzen veröffentlicht werden, bestimmt die gewählte Lizenz, ob ein Text als Trainingsmaterial genutzt werden darf. Bei den offeneren Lizenzen CC BY und CC BY SA ist dies möglich. Non-Commercial CC Lizenzen ermöglichen die Nutzung eines Textes zu Trainingszwecken, wenn alle Aspekte des Trainingsprozesses sowie Verfügbarmachung des Modells keinen kommerziellen Zwecken dient. Material mit Non-Derivatives CC Lizenz darf nicht zu Trainingszwecken verwendet werden (siehe [Creative Commons: Using CC-Licensed Works for AI Training](#)). In der Praxis kommt es jedoch immer wieder zur Nutzung von urheberrechtlich geschütztem Material (auch aus Raubkopien, siehe [Frage 22](#)) zu Trainingszwecken (siehe z.B. [Reisner 2025](#)).

24. Wie kann ich Konflikten, die aus KI-Nutzung resultieren könnten, vorbeugen?

Besonders im jetzigen Stadium, in dem die Arbeit mit generativer KI bei vielen Forschenden noch nicht vollständig etabliert ist und Empfehlungen/Richtlinien noch im Entstehen sind, sollten Forschende sich zu Beginn mit den für sie maßgeblichen KI-Richtlinien auseinandersetzen sowie ihre eigene Expertise in Bezug auf KI kritisch reflektieren. Für Promotionen sollten – sofern es von Institutsseite keine Richtlinien gibt – zwischen Betreuungspersonen und Doktorand:innen gemeinsame Regelungen abgesprochen werden. Arbeiten Forschende im Team oder einem langfristigen Forschungsprojekt mit fluktuierendem Personal (z.B. Editionsprojekte), sollte es stets eine offene Kommunikation über die Nutzung von KI sowie eine transparente interne Dokumentation darüber geben. Für andere Forschende im Team/Projekt sollte ebenso ersichtlich sein, ob und wie KI zum Einsatz kam. Dies gilt auch für andere Produkte wissenschaftlicher Arbeit abseits von „klassischen“ Publikationen, z.B. Software, Vortragsmanuskripte und Folien (die ggf. geteilt werden), Vorlagen für Posterpräsentationen und Weiteres.

25. Welche KI-Richtlinien sind für mich maßgeblich?

Dies hängt u.a. davon ab, wozu KI-Anwendungen in welchem Kontext genutzt werden. Für die Erstellung von Publikationen sollten neben institutionellen Richtlinien – sofern vorhanden – die KI-Policies der

Verlage beachtet werden. Existieren Richtlinien innerhalb der eigenen Fachcommunity, sollten diese Anwendung finden, insbesondere wenn diese strengere Kriterien beinhalten. Bei Promotionen gelten Promotions- und Prüfungsordnungen. Außerdem kann und sollte das Thema mit den Betreuungspersonen besprochen werden. Bei Anträgen gelten die Angaben aus [Frage 12](#). Sollten sich Widersprüche auftun – z.B. zwischen institutionellen Richtlinien und Verlagsrichtlinien – sollte diese frühzeitig kommuniziert werden.

26. Warum gibt es keine klaren Empfehlungen vom Ombudsgremium für die wissenschaftliche Integrität in Deutschland zum Thema KI und GWP?

Das Ombudsgremium für die wissenschaftliche Integrität in Deutschland (OWID) hat sich Anfang 2023 (damals noch unter dem Namen Ombudsman für die Wissenschaft) mit der Frage beschäftigt, ob zum Thema KI und GWP Empfehlungen herausgegeben werden sollten. In diesem Zusammenhang konzipierten und organisierten die Mitarbeiter:innen des Projektes „Dialogforen zur Stärkung der guten wissenschaftlichen Praxis“, Katrin Frisch, Felix Hagenström und Nele Reeg, einen Expert:innenworkshop, in dem die wichtigsten Fragen zu KI entlang der bestehenden GWP-Leitlinien des DFG-Kodex diskutiert wurden. Die Ergebnisse dieses Workshops mündeten in einem [Kurzbericht](#) sowie [einer Publikation zum Thema](#), die im Dezember 2023 in *der Zeitschrift für Bibliothekswesen und Bibliographie* (ZfBB) erschien. Da sich während des Workshops herauskristallisierte, dass viele Fragen zur GWP-konformen Nutzung von KI entweder schon durch den Kodex und weiteren Richtlinien adressiert werden oder aber einer Klärung in den einzelnen Fachcommunitys bedürfen, um verschiedener Fachspezifika gerecht zu werden, wurden keine allgemeinen Empfehlungen erarbeitet. OWID kann aber bei der Aufsetzung von Empfehlungen beraten.

Referenzen

Alle Referenzen finden sich als Direktlink ebenso an den jeweiligen Stellen im Text, an denen auf sie verwiesen wird. Eine [Literatursammlung zu KI und GWP](#) findet sich auch auf der Webseite des OWID.

Literatur

Bahr, Amrei (2024): Bridges Without Foundation? Why the Use of AI Tools in Academia Needs to Build on Ethics First. *FILozOFIA*, 79, No 5, pp. 486 – 500. <https://doi.org/10.31577/filozofia.2024.79.5.2>

Bergstrom, Carl T. und Joe Bak-Coleman (2025). AI, peer review and the human activity of science. *Nature*. <https://doi.org/10.1038/d41586-025-01839-w>

Cooper, A.F., Gokaslan, A., Cyphert, A.B., Sa, C.D., Lemley, M.A., Ho, D.E., & Liang, P. (2025). Extracting memorized pieces of (copyrighted) books from open-weight language models. ArXiv. <https://doi.org/10.48550/arXiv.2505.12546>

Frisch, Katrin, Felix Hagenström und Nele Reeg (2023). Textgenerierende KI und gute wissenschaftliche Praxis. *Zeitschrift für Bibliothekswesen und Bibliographie* 70, Nr. 6: 326–36. <https://doi.org/10.3196/186429502370667>.

Gao, Catherine A., Frederick M. Howard, Nikolay S. Markov et al. (2023). Comparing scientific abstracts generated by ChatGPT to real abstracts with detectors and blinded human reviewers, *npj Digital Medicine*, Volume 6, Issue 75. <https://doi.org/10.1038/s41746-023-00819-6>

Hao, Karen (2025). *Empire of AI: Dreams and Nightmares in Sam Altman's OpenAI*. New York: Penguin Press.

He, Jessica, Stephanie Houde, und Justin D. Weisz (2025). Which Contributions Deserve Credit? Perceptions of Attribution in Human-AI Co-Creation. *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems (CHI '25)*. 540, 1–18. <https://doi.org/10.1145/3706598.3713522>

Henderson, Peter, Xuechen Li, Dan Jurafsky, et al. (2023). Foundation Models and Fair Use. Arxiv. <https://doi.org/10.48550/arXiv.2303.15715>

Hinchliff Pearson, Sarah (2025). Understanding CC Licenses and AI Training: A Legal Primer. <https://creativecommons.org/2025/05/15/understanding-cc-licenses-and-ai-training-a-legal-primer/>

Hosseini, Mohammad, Gordijn, Bert, Kaebnick, Gregory E., & Holmes, Katie (2025). Disclosing generative AI use for writing assistance should be voluntary. *Research Ethics*. <https://doi.org/10.1177/17470161251345499>

Hosseini, Mohammad, David B. Resnik und Katie Holmes (2023). The ethics of disclosing the use of artificial intelligence tools in writing scholarly manuscripts. *Research Ethics*. Volume 19, Issue 4. <https://doi:10.1177/17470161231180449>

Hosseini, Mohammad, Serge P.J.M. Horbach (2023). Fighting reviewer fatigue or amplifying bias? Considerations and recommendations for use of ChatGPT and other large language models in scholarly peer review. *Res Integr Peer Rev* 8, 7. <https://doi.org/10.1186/s41073-023-00133-5>

Konertz, Roman (2023). Urheberrechtliche Fragen der Textgenerierung durch Künstliche Intelligenz: Insbesondere Schöpfungen und Rechtsverletzungen durch GPT und ChatGPT. *Wettbewerb in Recht und Praxis*, Jahrgang 69. <https://doi.org/10.18445/20230613-111651-0>

Messeri, Lisa, Crockett, M.J. (2024). Artificial intelligence and illusions of understanding in scientific research. *Nature* 627, 49–58. <https://doi.org/10.1038/s41586-024-07146-0>

Oertner, Monika (2024). ChatGPT als Recherchetool?: Fehlertypologie, technische Ursachenanalyse und hochschuldidaktische Implikationen. *Bibliotheksdienst* 58, Nr. 5, 259-297.

<https://doi.org/10.1515/bd-2024-0042>

Pope, Audrey (2024). NYT v. OpenAI: The Times's About-Face. *Harvard Law Review*.

<https://harvardlawreview.org/blog/2024/04/nyt-v-openai-the-times-about-face/>

Reisner, Alex (2025). The Unbelievable Scale of AI's Pirated-Books Problem. *The Atlantic*.

<https://www.theatlantic.com/technology/archive/2025/03/libgen-meta-openai/682093/>

Resnik, David B., Hosseini, Mohammad (2024). The ethics of using artificial intelligence in scientific research: new guidance needed for a new tool. *AI Ethics*. <https://doi.org/10.1007/s43681-024-00493-8>

Seadle, Michael (2017): *Quantifying Research Integrity*. San Rafael, California: Morgan & Claypool.

Schütze, Paul (2025). Zwischen Technologie und Ideologie: Ein Blick auf die un/realen Räume hinter KI. In: Clara Kindler-Mathôt/Didem Leblebici/Giacomo Marinsalta/Till Rückwart und Anna Zaglyadnova. *Un/Reale Interaktionsräume*. Bielefeld: transcript Verlag: 81-100.

<https://doi.org/10.14361/9783839471463-006>

Slattery, P., Saeri, A. K., Grundy, E. A. C., Graham, J., Noetel, M., Uuk, R., Dao, J., Pour, S., Casper, S., & Thompson, N. (2024). The AI Risk Repository: A Comprehensive Meta-Review, Database, and Taxonomy of Risks from Artificial Intelligence. <https://doi.org/10.48550/arXiv.2408.12622>

Verch, Ulrike (2024). Per Prompt zum Plagiat? Rechtssicheres Publizieren von KI-generierten Inhalten. *API*, Band 5, Nr. 1. <https://doi.org/10.15460/apimagazin.2024.5.1>.

Weaver, Kari D. (2024). The artificial intelligence disclosure (AID) framework: An introduction. *C&RL News*, 85(10), 407-411. <https://crln.acrl.org/index.php/crlnews/article/view/26548>

Weber-Wulff, Debora, Alla Anohina-Naumeca, Sonja Bjelobaba, et al. (2023): Testing of detection tools for AI-generated text. *International Journal for Education Integrity*, Volume 19, Issue 26.

<https://doi.org/10.1007/s40979-023-00146-z>

Weinberg, Michael (2025). Are AI Bots Knocking Cultural Heritage Offline? The GLAM-E Lab.

<https://www.glamelab.org/products/are-ai-bots-knocking-cultural-heritage-offline/>

Wolkovich, E.M. (2024). 'Obviously ChatGPT' — how reviewers accused me of scientific fraud. *Nature*; Career Column. <https://doi.org/10.1038/d41586-024-00349-5>

Policies, Handreichungen, Empfehlungen

ACS Nano: Best Practices for Using AI When Writing Scientific Manuscripts, 2023, 17, 5, 4091–4093. <https://doi.org/10.1021/acsnano.3c01544>

AI Attribution Toolkit: <https://aiattribution.github.io/>

American Psychological Association: How to Cite ChatGPT: <https://apastyle.apa.org/blog/how-to-cite-chatgpt>

Berlin Universities Publishing: „Leitlinie zum Umgang mit künstlicher Intelligenz“ und „Handreichung zur Zitation von KI-Tools“, 2024 <https://www.berlin-universities-publishing.de/ueber-uns/policies/ki-leitlinie/index.html>

Börsenverein des Deutschen Buchhandels: Künstliche Intelligenz im Verlagsbereich: FAQ zu generativer KI. <https://www.boersenverein.de/beratung-service/recht/kuenstliche-intelligenz/>

Chicago Manual of Style: Citing content developed or generated by artificial intelligence:
<https://www.chicagomanualofstyle.org/qanda/data/faq/topics/Documentation/faq0422.html>

Deutsche Forschungsgemeinschaft (DFG): Leitlinien zur Sicherung guter wissenschaftlicher Praxis, 2019: <https://wissenschaftliche-integritaet.de/>

DFG: Stellungnahme des Präsidiums der Deutschen Forschungsgemeinschaft (DFG) zum Einfluss generativer Modelle für die Text- und Bilderstellung auf die Wissenschaften und das Förderhandeln der DFG, 2023: <https://wissenschaftliche-integritaet.de/verwendung-generativer-modelle/>

Elsevier: Generative AI policies for journals: <https://www.elsevier.com/de-de/about/policies-and-standards/generative-ai-policies-for-journals>

EU AI Act (Gesetz über künstliche Intelligenz (Verordnung (EU) 2024/1689):
<https://artificialintelligenceact.eu/de/das-gesetz/>

Europäische Kommission: Living guidelines on the responsible use of generative AI in research, V2 2025: https://research-and-innovation.ec.europa.eu/document/download/2b6cf7e5-36ac-41cb-aab5-0d32050143dc_en?filename=ec_rtd_ai-guidelines.pdf

Frontiers: Artificial intelligence: fair use and disclosure policy:
<https://www.frontiersin.org/guidelines/policies-and-publication-ethics>

ICMJE: Recommendations for the Conduct, Reporting, Editing, and Publication of Scholarly Work in Medical Journals, 2024: <https://www.icmje.org/icmje-recommendations.pdf>

Modern Language Association: How do I cite generative AI in MLA style?: <https://style.mla.org/citing-generative-ai/>

PLOS One: Ethical Publishing Practice/ Artificial Intelligence Tools and Technologies:
<https://journals.plos.org/plosone/s/ethical-publishing-practice#loc-artificial-intelligence-tools-and-technologies>

Science Journals: Editorial Policies: <https://www.science.org/content/page/science-journals-editorial-policies#image-and-text-integrity>

Springer Nature: Editorial Policies: Artificial Intelligence:
<https://www.springernature.com/gp/policies/editorial-policies>

Taylor & Francis: AI Policy: <https://taylorandfrancis.com/our-policies/ai-policy/>

WAME: Recommendations on Chatbots and Generative Artificial Intelligence in Relation to Scholarly Publications, 2023: <https://wame.org/page3.php?id=106>

Wiley: AI Guidelines [for book authors]: <https://www.wiley.com/en-us/publish/book/resources/ai-guidelines/>

Wiley: Best Practice Guidelines on Research Integrity and Publishing Ethics: Artificial Intelligence:
<https://authorservices.wiley.com/ethics-guidelines/index.html#22>