

GiL to go. Lese- und Hörimpulse für eine geschlechtergerechte Lehre

Modul: Verwendung digitaler Tools

Arbeiten mit ChatGPT und Co

Stand: November 2025



Was hat ChatGPT mit Geschlechtergerechtigkeit zu tun?

Der Einsatz von Anwendungen großer Sprachmodelle (LLMs¹) wie zum Beispiel ChatGPT in der universitären Lehre ist ein viel diskutiertes Thema². Zum einen, weil LLMs möglicherweise unerkannt bei Studien- oder Prüfungsleistungen eingesetzt werden können und zum anderen, weil – Stand 2025 – in vielen Fachkulturen bisher nur sehr wenige sinnvolle und gut erprobte Einsatzmöglichkeiten für die universitäre Lehre zur Verfügung stehen.

Bezogen auf die Geschlechtergerechtigkeit in Lehrveranstaltungen kommt hinzu, dass LLMs nicht nur Geschlechtsstereotype reproduzieren, sondern auch ungewollt diskriminierende, queerfeindliche oder rassistische Äußerungen produzieren können (UNESCO 2024). Zudem besteht die Gefahr, dass es bei der Nutzung von LLMs geschlechtsbedingte ungleiche Voraussetzungen gibt.

¹ LLMs steht für Large Language Models (große Sprachmodelle), also generative KI-Anwendungen für die Produktion geschriebener Texte.

² Angehörige der Universität können das Bielefelder KI-Portal (BIKI) nutzen, welches möglichst herstellerunabhängig und datenschutzkonform gestaltet ist.





Fakten

LLMs lernen auf der Basis riesiger Textmengen, das jeweils folgende Wort vorherzusagen (Handschuh 2024). Sie errechnen anhand der Abfolge von Wörtern und Wortfragmenten, wie der Text weitergehen oder wie die Antwort auf eine Anfrage aussehen könnte. LLMs können so beispielsweise sehr schnell in natürlichen und formalen Sprachen verfasste Texte erstellen. Diese sind nicht sicher von solchen unterscheidbar, die Menschen verfasst haben ([Universität Bielefeld 2024](#)).

Was dabei an der Oberfläche nicht sichtbar wird, sind der große Rechenaufwand und der damit verbundene erhebliche Energieverbrauch. „Eine Anfrage an ChatGPT [verbraucht beispielsweise] etwa das 10- bis 300-fache an Strom einer Google Suche“ (Summer u. a. 2025, 43).

Geschlechtergerechtigkeit und Diskriminierung

LLMs spiegeln ohne zusätzliche Intervention alle statistischen Zusammenhänge aus den Trainingsdaten sowie den Anweisungen ihrer Programmierer*innen entsprechend wider.

Diskriminierende oder falsche Aussagen von LLMs

Die Modelle haben keinen eigenständigen werteorientierten Anspruch und unterscheiden nicht zwischen inhaltlich oder ethisch „richtigen“ und „falschen“ Aussagen. Da die Ergebnisse auf statistischen Auswertungen beruhen, differenzieren sie nicht explizit zwischen Inhalten, für die es gute Belege gibt und solchen, die nicht belegbar sind. Bei durch LLMs generierten Texten ist darüber hinaus häufig nicht erkennbar, wer die ursprünglichen Urheber*innen der ausgegebenen Inhalte sind und

2



Dieses Werk ist freigegeben unter der [Creative-Commons-Lizenz CC BY-NC 4.0](#) (Weitergabe unter gleichen Bedingungen). Diese Lizenz gilt nur für das Originalmaterial. Alle gekennzeichneten Fremdinhalte (z.B. Abbildungen, Fotos, Tabellen, Zitate etc.) sind von der CC-Lizenz ausgenommen. Für deren Wiederverwendung ist es ggf. erforderlich, weitere Nutzungsgenehmigungen beim jeweiligen Rechteinhaber einzuholen.

wie seriös die Informationsquelle ist. Es könnte sich dabei gleichermaßen um einen wissenschaftlichen Aufsatz oder schlichtweg um Fake News handeln. Außerdem können Aussagen generiert werden, die mit den ethischen Ansprüchen einer demokratischen Gesellschaft nicht übereinstimmen. Sind in den Trainingsdaten zum Beispiel stereotype Rollenvorstellungen oder diskriminierende Äußerungen enthalten, so werden diese von der KI zunächst unkritisch übernommen. Eine Studie der UNESCO (Deutsche UNESCO Kommission 2023) zeigt, dass LLMs dazu tendieren, Geschlechtsstereotype zu reproduzieren und diese somit zu verstärken. Die generierten Texte können daher sexistische, queerphobe oder rassistische Inhalte enthalten.

Ein unreflektierter Umgang mit LLMs kann demnach zu einer Reproduktion und Verfestigung solcher Ansichten führen.

Reproduktion von geschlechtsbezogenen Stereotypen

KI-Anwendungen tendieren bei geschlechtsneutralen Suchbegriffen dazu, geschlechtsstereotype Bild- und Textergebnisse zu produzieren (Franken/Mauritz 2020: 8; Marsden 2023). Gibt man zum Beispiel in einem Prompt „Person, die im Tischlerhandwerk arbeitet“ oder Tischler*in als geschlechtsneutralen Begriff für diese Berufsgruppe ein, werden wahrscheinlich ausschließlich männliche Personen erwähnt. Um die Namen weiblicher Personen zu erfahren, muss unter Umständen explizit die weiblicher Form „Tischlerin“ gepromptet werden. Bei der Anfrage nach pflegenden Personen verhält es sich anders herum. Hier werden eher weibliche Personen genannt.



Gleichberechtigte Nutzung von LLMs?

Weiterhin ist für einen geschlechtergerechten Einsatz von LLMs bedeutsam, dass männliche Personen KI-Anwendungen häufiger und unbefangener nutzen als weibliche (Armutat u. a. 2024). Laut Studien nehmen Männer KI-Technologien insgesamt positiver wahr als Frauen und befürworten generell deutlich häufiger deren Nutzung und die Entwicklung (Armutat u. a. 2024). Zu TIN*(also trans, inter und nicht-binäre) Personen konnten keine aussagekräftigen Studien gefunden werden.

Bei der Anwendung von KI-Technologien schreiben sich Männer häufig mehr Kompetenz zu, auch wenn dies nicht notwendigerweise einen Rückschluss auf tatsächliche geschlechtsbedingte Kompetenzunterschiede gibt (Franken/Mauritz 2020: 4ff).

Frauen hingegen schätzen sich selbst in Bezug auf LLMs häufig als weniger kompetent ein (Armutat u. a. 2024). Dies kann in einer zögerlicheren Nutzung resultieren und dadurch unter Umständen zu ungleichen Chancen führen. Ein kompetenter Umgang mit LLMs kann in bestimmten Bereichen auch für die Vorbereitung auf einen zukünftigen Arbeitsmarkt von großer Bedeutung sein.

Über-, Unter- bzw. Nicht-Repräsentation sozialer Gruppen

Auch eine Über-, Unter- bzw. Nicht-Repräsentation sozialer Gruppen im Internet wird durch die LLMs wiedergegeben. Während die Ansichten einiger Personengruppen überproportional vertreten sind, werden die Ansichten und Einstellungen bestimmter Personen(-gruppen) nur wenig im Internet abgebildet. Diese Personen haben beispielsweise aus kulturellen oder infrastrukturellen Gründen oder aufgrund von Sorgearbeit oder Beeinträchtigungen wenig Möglichkeiten, im digitalen Raum sichtbar zu werden oder schlicht sehr wenig Zeit. Folglich sind ihre Themen und



Ansichten auch in den generierten Texten entsprechend selten, oder eben nur aus der Außenperspektive aus Sicht der überrepräsentierten Gruppen, vertreten.

Bei der Frage nach einem Burnout wird zum Beispiel als Ursache vor allem Überforderung durch Erwerbsarbeit aufgeführt. Dass auch Personen, die Sorgearbeit leisten, ein Burnout erleiden können, wird unter Umständen nicht erwähnt. Im privaten Bereich werden hier eher Ursachen wie Perfektionismus oder eine ungenügende Abgrenzungsfähigkeit aufgeführt.

Unterstützung bei der Erstellung gendersensibler Inhalte

Auf der anderen Seite können LLMs und andere KI-Anwendungen mit einer klugen Auswahl von Prompts³ explizit dabei unterstützen, Lehrinhalte und Arbeitsergebnisse diskriminierungsarm zu gestalten.



Was kann ich tun?

Nutzung von LLMs für eigene (geschlechtergerechte) Inhalte

Prompts entwickeln

Es erleichtert die Arbeit mit LLM-generierten Texten ungemein, wenn schon die Prompts sorgfältig geplant sind. Hier kann es helfen, klare und explizite Anweisungen zu geben. Wenn zum Beispiel geschlechtergerechte Sprache verwendet werden soll, kann dies direkt im Prompt eingebunden werden. Das LLM kann auch den Auftrag bekommen, ausschließlich wissenschaftliche Texte oder nur ganz bestimmte Texte zu nutzen. Vorabtests der Prompts können die Wahrscheinlichkeit erhöhen, dass sie das

³ Als Prompt wird die Eingabeaufforderung für ein LLM bezeichnet.



gewünschte Ergebnis liefern. Allerdings bleibt auch hier die Gefahr, dass die produzierten Texte fehlerhaft sind.

Texte prüfen

Es ist unbedingt notwendig, die von einem LLM generierten Texte immer sorgfältig zu überprüfen, bevor sie weiterverwendet werden.

Integration von LLMs für Studierende in Lehrveranstaltungen

Folgende Aspekte können einen guten Umgang mit LLMs in der Lehre fördern.

Transparenz:

Transparenz bezüglich der Verwendung von LLMs kann helfen, Unsicherheiten zu vermeiden. Es ist zum Beispiel wichtig schon zu Beginn einer Veranstaltung, offen darzulegen, wie in Prüfungssituationen und bei Aufgaben im Kurs mit dem Einsatz von LLMs verfahren werden soll. Falls LLMs in der Veranstaltung explizit einbezogen werden, kann ebenfalls dargelegt werden, welche Möglichkeiten es gibt, sich in die Nutzung einzuarbeiten.

Bewusstsein schaffen:

Den Umgang mit diskriminierenden oder sexistischen oder geschlechterstereotypen Ergebnissen von LLMs zu thematisieren und für deren Erkennung zu sensibilisieren, kann helfen, Biases (also systematische Fehleinschätzungen) zu vermeiden.

Wissen über die Chancen und Risiken von LLMs auch im Kontext von Geschlechtergerechtigkeit, Diversity und Nachhaltigkeit kann Studierenden bei einem verantwortungsvollen Einsatz von LLMs helfen. Hier können auch interdisziplinäre Ansätze oder Perspektiven aus den Gender Studies förderlich sein. Da nicht jede Lehrperson solches Wissen vermitteln kann oder möchte, kann es hilfreich sein, auf die zahlreichen [Angebote des Zentrums für Lehren und Lernen](#) der



Universität Bielefeld oder die „[Didaktische Handreichung zur Nutzung generativer KI in Studium und Lehre](#)“ (Zentrum für Lehren und Lernen 2025) zurückgegriffen werden. Zudem kann die (kritische) Auseinandersetzung mit LLMs im Rahmen universitärer Lehre auch eine Möglichkeit sein, die Studierenden grundsätzlich an Themen der Geschlechtergerechtigkeit heranzuführen.

Raum für Experimente geben:

Eine weitere Möglichkeit, die Sensibilität für Biases zu erhöhen, ist das Experimentieren mit unterschiedlichen Prompts. Studierende können so ausprobieren, welche Ergebnisse diese hervorbringen und darüber ins Gespräch kommen. Dabei können beispielsweise solche Prompts verwendet werden, die mit hoher Wahrscheinlichkeit einen Gender Bias erzeugen (z.B. „10 beste Philosophen“ vs. „10 beste Philosoph*innen“) aber auch fachliche Fragen, die schlussendlich auf ihre Validität überprüft werden können.

Eigenes Wissen nutzen und erweitern:

Gute KI-Ergebnisse erfordern ein fundiertes Vorwissen und eine klare Vorstellung vom gewünschten Arbeitsergebnis. Übungen zum Erkennen sachlich falscher oder nicht geschlechtergerechter Aussagen können ebenfalls förderlich sein.

Partizipation fördern:

Wenn die Nutzung von LLMs in einem Kurs explizit gefordert wird oder erlaubt ist, sollten alle Studierenden gleichermaßen befähigt und dabei unterstützt werden, die Anwendungen effizient zu nutzen und mit den entsprechenden Arbeitsergebnissen verantwortungsvoll umzugehen. Gezielte Übungen in den Seminarsitzungen können helfen. Soll es zum Beispiel erlaubt sein, das LLM zur Unterstützung bei der Themenfindung, dem Entwickeln einer Fragestellung oder der Literaturrecherche zu



nutzen, können diese Schritte gemeinsam mit mehr oder weniger spezifischen Prompts ausprobiert und die Ergebnisse im Kurs gemeinsam ausgewertet werden.

Geschlechtergerechtigkeit als Aufgabe für LLMs:

Die Anwendung kann zum Beispiel die Aufgabe bekommen, vorhandene Inhalte auf eine durchgängige Anwendung geschlechtergerechter Sprache oder unbewusste Vorurteile hin zu analysieren. Eine diskriminierungsarme Gestaltung von Inhalten kann auch bei einer Text- oder Bildproduktion direkt eingefordert werden. Die Umsetzung solcher Aufgaben sollte aber auf jeden Fall abschließend noch einmal überprüft werden, denn es kann immer sein, dass die Modelle sie unvollständig erledigen oder die Aufgabe nicht im Sinne der auftraggebenden Person durchführen.



Literatur

Armutat, Sascha, Malte Wattenberg, und Nina Mauritz. 2024. „Gender Bias bei der Wahrnehmung Künstlicher Intelligenz: Ursachen und Strategien“. *Femina Politica – Zeitschrift für feministische Politikwissenschaft* 33 (2–2024): 125–28. <https://doi.org/10.3224/feminapolitica.v33i2.14>.

Deutsche UNESCO Kommission. 2023. *UNESCO-Empfehlung zur Ethik der Künstlichen Intelligenz*. Deutsche UNESCO-Kommission e.V.

Handschuh, Siegfried. 2024. „Grosse Sprachmodelle“. *Informationswissenschaft: Theorie, Methode und Praxis* 8 (1): 1. <https://doi.org/10.18755/iw.2024.3>.

Summer, Alexander, Marcus Bentele, und Sabrina Schneider. 2025. *Künstliche Intelligenz im Spannungsfeld von Nachhaltigkeit und Ressourcenverbrauch*. Klima: Verhalten? Gestalten!: Beiträge zum uDay XXIII-25. Juni 2025: 2847 KB,



Dieses Werk ist freigegeben unter der [Creative-Commons-Lizenz CC BY-NC 4.0](https://creativecommons.org/licenses/by-nc-sa/4.0/) (Weitergabe unter gleichen Bedingungen). Diese Lizenz gilt nur für das Originalmaterial. Alle gekennzeichneten Fremdinhalte (z.B. Abbildungen, Fotos, Tabellen, Zitate etc.) sind von der CC-Lizenz ausgenommen. Für deren Wiederverwendung ist es ggf. erforderlich, weitere Nutzungsgenehmigungen beim jeweiligen Rechteinhaber einzuholen.

69 pages. Application/pdf, 2847 KB, 69 pages. <https://doi.org/10.25924/OPUS-6980>.

UNESCO. 2024. „Challenging Systematic Prejudices: An Investigation into Bias against Women and Girls in Large Language Models - UNESCO Digital Library“. International Research Centre on Artificial Intelligence under the auspices of UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000388971>.

Universität Bielefeld. 2024. „Frag BIKI - Was bist Du, BIKI?“ Blog. Juni 19. <https://blogs.uni-bielefeld.de/blog/biki/entry/was-bist-du-biki#mainSection>.

Zentrum für Lehren und Lernen, Universität Bielefeld. 2025. „Didaktische Handreichung zur Nutzung generativer KI in Studium und Lehre“. <https://www.uni-bielefeld.de/lehre/digitale-lehre/ki-tools/2025-09-Didaktische-Handreichung-genKI.pdf>.

Text: Dr. Beate Lingnau, Nora Charlotte Gehlen, Lisa Steven

Beratung: Benjamin Angerer (ZLL)

Beitrag: Sophie Marticke, Ragna Verhoeven

GiL to go ist eine Produktion des Gleichstellungsbüros der Universität Bielefeld aus dem Programm Gleichstellung sehen und hören. www.uni-bielefeld.de/gleichstellung-sehen-und-hoeren



Dieses Werk ist freigegeben unter der [Creative-Commons-Lizenz CC BY-NC 4.0](https://creativecommons.org/licenses/by-nc-sa/4.0/) (Weitergabe unter gleichen Bedingungen). Diese Lizenz gilt nur für das Originalmaterial. Alle gekennzeichneten Fremdinhalte (z.B. Abbildungen, Fotos, Tabellen, Zitate etc.) sind von der CC-Lizenz ausgenommen. Für deren Wiederverwendung ist es ggf. erforderlich, weitere Nutzungsgenehmigungen beim jeweiligen Rechteinhaber einzuholen.

Konzept: Dr. Beate Lingnau, Nora Charlotte Gehlen, Siân Birkner

Redaktion: Sandra Sieraad

Gestaltung: © Universität Bielefeld | Referat für Kommunikation, Grafik | 2025

Barrierefreiheit: Marcel Gemander (ZAB)



Dieses Werk ist freigegeben unter der [Creative-Commons-Lizenz CC BY-NC 4.0](https://creativecommons.org/licenses/by-nc-sa/4.0/) (Weitergabe unter gleichen Bedingungen). Diese Lizenz gilt nur für das Originalmaterial. Alle gekennzeichneten Fremdinhalte (z.B. Abbildungen, Fotos, Tabellen, Zitate etc.) sind von der CC-Lizenz ausgenommen. Für deren Wiederverwendung ist es ggf. erforderlich, weitere Nutzungsgenehmigungen beim jeweiligen Rechteinhaber einzuholen.