

KI-generierte Inhalte erkennen: Verdachts- und Plausibilitätsprüfung

Dieser Abschnitt unterstützt Lehrende, Mitarbeitende und Betreuer*innen bei der **sachlichen Prüfung von Verdachtsfällen**, wenn bei studentischen Arbeiten Anhaltspunkte für **nicht offengelegte, unzulässige oder irreführend dargestellte KI-Nutzung** bestehen.

Im Mittelpunkt steht nicht der technische „Nachweis von KI“, da dieser nach aktuellem Stand nicht zuverlässig möglich ist, sondern eine **dokumentierte Gesamtwürdigung** von Eigenleistung, Transparenz, Quellen, Daten, Abbildungen, Argumentation und Nachvollziehbarkeit des Arbeitsprozesses.

Die Anleitung beschreibt einen praxisnahen Prüfworkflow von der Klärung der erlaubten KI-Nutzung über die inhaltliche Prüfung und das Anfordern von Prozessnachweisen bis hin zum Gespräch mit Studierenden und zur dokumentierten Entscheidung.

Sie enthält außerdem Warnsignale für KI-generierte Texte, Daten, Tabellen, Abbildungen und Quellen sowie Hinweise zum verantwortungsvollen Umgang mit KI-Detektoren, Perplexität, Burstiness und möglichen falschpositiven Ergebnissen.

- [Entscheidungshilfe für Mitarbeitende, Lehrende und Betreuer*innen](#)
- [Hilfestellung zur Textanalyse](#)
- [Prüfmethoden für Daten, Tabellen und statistischen Auswertungen](#)
- [Prüfmethoden für Abbildungen, Diagrammen und Bildmaterial](#)
- [Übersicht: Merkmale KI-generierter Inhalte](#)

Entscheidungshilfe für Mitarbeitende, Lehrende und Betreuer*innen

Ausgangsfall

Bei einer schriftlichen Arbeit, Abschlussarbeit, Forschungsarbeit, Seminararbeit, Hausarbeit, Fallausarbeitung, einem Laborprotokoll oder einer sonstigen prüfungsrelevanten Leistung bestehen Anhaltspunkte dafür, dass KI-generierte Texte, KI-generierte Daten, KI-generierte Abbildungen oder sonstige KI-generierte Inhalte nicht offengelegt, unzulässig verwendet oder als eigene Leistung ausgegeben wurden.

Grundsatz

Die Prüfung dient nicht dem „Nachweis von KI“ im technischen Sinn (diese ist nach aktuellem Stand nicht 100% möglich), sondern der **begründeten Verdachtsprüfung und Plausibilitätsprüfung**.

Maßgeblich ist eine dokumentierte Gesamtwürdigung:

- In welcher Kategorie (A – D) ist die Arbeit erstellt worden? → KI-Richtlinie § 5
- Ist die Eigenleistung erkennbar?
- Wurde eine erlaubte KI-Nutzung transparent dokumentiert?
- Sind Text, Daten, Quellen, Abbildungen und Argumentation fachlich plausibel?
- Kann die oder der Studierende die Arbeitsschritte nachvollziehbar erklären?
- Wurden Vorgaben der Lehrveranstaltung, der Betreuung oder der KI-Richtlinie eingehalten?

Ein bloßer Eindruck, ein ungewöhnlicher Schreibstil oder ein einzelnes Verdachtsmoment reichen nicht für eine belastbare Entscheidung.

Empfehlung

Besonders wirksam ist nicht die Textanalyse allein, sondern die Nachvollziehbarkeit der Eigenleistung. Es ist daher empfehlenswert folgende Elemente in die Aufgabenstellung einzubauen: ein Forschungstagebuch, dokumentierte Zwischenstände, Vor-Versionen, ein Quellenprotokoll,

Abgabe der Rohdaten oder eine mündliche Vorstellung der Arbeit.

Empfohlener Prüfworkflow

Stufe 1: Ausgangslage klären

Was war in der Lehrveranstaltung erlaubt? Gab es Kennzeichnungspflicht? Wurde eine Eigenständigkeitserklärung verlangt? Wurden KI-Nutzungsregeln schriftlich kommuniziert?

Stufe 2: Arbeit inhaltlich prüfen

Quellen, Argumentation, Methode, Daten, Abbildungen und Schlussfolgerungen prüfen. Besondere Anzeichen sind erfundene Quellen, nicht reproduzierbare Daten, widersprüchliche Tabellen, perfekte Statistik, fehlende Rohdaten oder ein gut klingender, aber inhaltsleerer Text.

Stufe 3: Prozessnachweise von Studierenden anfordern

Zwischenversionen, Notizen, Literaturprotokoll, Datenfile, Analysecode, Prompt-/KI-Nutzungsdokumentation, Bildrohdateien oder Entwurfsstände.

Stufe 4: Gespräch mit Studierenden

Studierende sollen zentrale Argumente, Quellenwahl, Methode, Datenaufbereitung und Abbildungen erklären. Bei Verdacht kann eine kurze mündliche Verteidigung oder Nachberechnung einzelner Ergebnisse verlangt werden.

Stufe 5: Dokumentierte Gesamtwürdigung

Entscheidung nur auf Basis mehrerer Indizien: Regelverstoß, fehlende Kennzeichnung, fehlende Eigenleistung, nicht erklärbare Inhalte, nicht reproduzierbare Daten oder nachweislich erfundene Quellen/Abbildungen

Hilfestellung zur Textanalyse

Da KI und menschliche Texte überlappen, gibt es aktuell keine 100% technische Methode KI-generierte Texte zu identifizieren. Die meisten KI-Detektoren basieren auf dem Prinzip der Perplexität. Niedrige Perplexität (KI-Verdacht): Der Text wirkt glatt, jedes Wort folgt logisch und erwartbar auf das vorherige, ohne syntaktische "Stolpersteine". Hohe Perplexität (Mensch-Indiz): Der Text enthält Brüche, seltene Wortwahl oder statistisch unwahrscheinliche Satzkonstruktionen.

Neben der Perplexität wird oft die "Burstiness" herangezogen, die den Rhythmus und die Variation der Satzlänge betrachtet. Menschen wechseln oft unregelmäßig zwischen kurzen und langen, verschachtelten Sätzen, während KI-Modelle tendenziell eine monotone, gleichmäßige Satzstruktur produzieren.

Da auch Menschen erwartbaren Text schreiben, liegt die Rate der falschpositiven KI-Zuweisungen bei etwa $\leq 20\%$. Diese Rate kann über 60% betragen, wenn Menschen nicht in ihrer Muttersprache schreiben (McKie A., 2026, Nature 655: 535-537). Hybride Texte mit KI und menschlichen Anteilen oder KI-Texte die „humanisiert“ wurden entgehen häufig der Detektion.

Hier ist eine Liste mögliche Auffälligkeiten in KI-generierten Texten

- deutlicher Stilbruch gegenüber früheren Arbeiten oder bekannten Schreibproben
- auffällig generische, glatte oder inhaltlich unpräzise Sprache
- Quellen, die nicht auffindbar sind oder nicht zur Aussage passen
- Methodenbeschreibung, Daten und Ergebnisse passen nicht zusammen
- Tabellen, Prozentwerte, Fallzahlen oder statistische Angaben sind widersprüchlich
- Abbildungen enthalten sachliche, anatomische, technische oder grafische Unplausibilitäten
- KI-Nutzung wurde nicht oder nur unvollständig dokumentiert
- Studierende können zentrale Inhalte, Quellen, Daten oder Arbeitsschritte nicht erklären
- Arbeit enthält Spuren von Prompts, KI-Formulierungen, KI-Formatierungen, Platzhaltern oder generischen Textbausteinen
- Im Anhang finden Sie weiter Hinweise auf KI-generierte Texte

Hier ist eine umfangreiche Liste mit Warnsignalen, die auf KI-Nutzung hindeuten

Stil und Sprache

Prüfen:

- Wirkt der Text ungewöhnlich glatt, allgemein oder austauschbar?
- Gibt es starke Stilbrüche zwischen einzelnen Kapiteln?
- Weicht der Stil deutlich von früheren Arbeiten, Exposés, E-Mails oder Zwischenfassungen ab?

- Enthält der Text viele richtige, aber oberflächliche Aussagen ohne eigene fachliche Vertiefung?
- Werden Fachbegriffe formal korrekt, aber kontextuell ungenau verwendet?
- Sind Übergänge und Zusammenfassungen sehr generisch?
- Fehlen individuelle Begründungen, lokale Bezüge, Kursbezüge oder konkrete Forschungsentscheidungen?

Bewertung:

- Sprachliche Glätte allein ist kein Beweis.
- Ein Stilbruch ist ein Hinweis, wenn er mit weiteren Auffälligkeiten zusammentrifft.
- Entscheidend ist, ob die fachliche Eigenleistung nachvollziehbar bleibt.

Argumentation und fachliche Tiefe

Prüfen:

- Wird eine klare eigene Fragestellung entwickelt?
- Werden Begriffe sauber definiert?
- Wird Literatur kritisch eingeordnet oder nur zusammengefasst?
- Passen Einleitung, Forschungsfrage, Methode, Ergebnisse und Schlussfolgerung zusammen?
- Werden Gegenargumente, Limitationen und Unsicherheiten fachlich reflektiert?
- Gibt es scheinbar überzeugende, aber inhaltlich leere Passagen?
- Werden Schlussfolgerungen gezogen, die aus den Ergebnissen nicht folgen?

Warnsignale:

- viele abstrakte Aussagen ohne konkrete Belege
- fehlende Verbindung zwischen Theorie und eigener Analyse
- unplausible Sicherheit bei unsicherer Datenlage
- scheinbar professionelle, aber fachlich ungenaue Argumentation
- fehlende Auseinandersetzung mit naheliegender Pflichtliteratur oder zentralen Quellen

Quellen und Zitate

Prüfen:

- Existieren die angegebenen Quellen tatsächlich?
- Stimmen Titel, Autor*innen, Jahr, DOI, Verlag, Zeitschrift und Seitenzahlen?
- Unterstützt die Quelle wirklich die behauptete Aussage?
- Sind Zitate korrekt wiedergegeben?
- Gibt es ungewöhnlich viele Quellen, die schwer auffindbar sind?
- Werden nicht einschlägige oder veraltete Quellen als Hauptbelege verwendet?
- Gibt es erfundene DOI, falsche Zeitschriftentitel oder unplausible Kombinationen von Autor*innen und Themen?

Warnsignale mit hoher Bedeutung:

- nicht auffindbare Quellen
- erfundene oder falsche DOI
- Quellen, die zwar existieren, aber die zitierte Aussage nicht enthalten
- systematische Fehlzitationen
- Literaturverzeichnis passt nicht zu den Fußnoten oder In-Text-Zitaten

Methodik und fachlicher Aufbau

Prüfen:

- Ist die Methode für die Forschungsfrage geeignet?
- Sind Stichprobe, Material, Datenerhebung und Auswertung nachvollziehbar beschrieben?
- Sind die verwendeten Fachbegriffe korrekt?
- Passen Methodenbeschreibung und Ergebnisdarstellung zusammen?
- Werden Limitationen realistisch dargestellt?
- Wurde eine Methode beschrieben, die tatsächlich durchgeführt worden sein kann?

Warnsignale:

- Methodenteil wirkt wie ein generischer Lehrbuchtext
- fehlende Angaben zu konkreten Durchführungsschritten
- Ergebnisse werden berichtet, ohne dass klar ist, wie sie entstanden sind
- verwendete Verfahren passen nicht zum Datentyp
- Schlussfolgerungen gehen deutlich über die Methode hinaus
- ethische, datenschutzrechtliche oder studienbezogene Voraussetzungen fehlen trotz sensibler Daten

Prüfmethoden für Daten, Tabellen und statistischen Auswertungen

Rohdaten und Dokumentation prüfen

Anfordern oder prüfen:

- Rohdatensatz
- Codebook oder Variablenerklärung
- Erhebungsinstrument
- Ein- und Ausschlusskriterien
- Datenerhebungszeitraum
- Dokumentation der Datenbereinigung
- Analysecode oder Auswertungsschritte
- Versionen von Tabellen und Auswertungen
- Ethikvotum, Datenschutzfreigabe oder sonstige Genehmigungen, soweit erforderlich

Konsistenzprüfung

Prüfen:

- Stimmen Fallzahlen in Text, Tabellen und Abbildungen überein?
- Passen Prozentwerte zu den angegebenen absoluten Zahlen?
- Sind Mittelwerte, Standardabweichungen, Konfidenzintervalle und p-Werte plausibel?
- Stimmen Tabellenüberschriften, Variablen und Legenden?
- Werden dieselben Gruppenbezeichnungen konsistent verwendet?
- Sind Rundungen nachvollziehbar?
- Gibt es Ergebnisse ohne vorher beschriebene Analyse?
- Werden Analysen beschrieben, die in den Ergebnissen nicht auftauchen?

Warnsignale:

- wechselnde Fallzahlen ohne Erklärung
- Prozentwerte, die sich nicht aus den Fallzahlen ergeben
- identische Werte an ungewöhnlich vielen Stellen
- zu perfekte Verteilungen
- fehlende Ausreißer bei realistisch streuenden Daten
- Korrelationen oder Effekte, die methodisch nicht plausibel sind
- Tabellen, die wie generische Musterbeispiele wirken

Reproduzierbarkeit prüfen

Prüfen:

- Kann die Auswertung anhand der Rohdaten nachvollzogen werden?
- Können zentrale Ergebnisse mit einfachen Nachrechnungen bestätigt werden?
- Gibt es eine nachvollziehbare Verbindung zwischen Rohdaten, Analyse und Ergebnisdarstellung?
- Sind Transformationen, Ausschlüsse und Bereinigungen dokumentiert?
- Können Studierende erklären, warum bestimmte statistische Verfahren gewählt wurden?

Mögliche Maßnahmen:

- einzelne Berechnungen gemeinsam durchgehen
- Studierende zentrale Ergebnisse nachrechnen lassen
- Auswertungsschritte erläutern lassen
- Analysecode oder Tabellenformeln prüfen
- Abgleich mit Betreuungsprotokollen, Laborbuch oder Forschungstagebuch

Spezifische Warnsignale bei KI-generierten oder erfundenen Daten

Warnsignale:

- Daten wirken vollständig, sauber und widerspruchsfrei, obwohl das Erhebungsdesign Ausfälle erwarten lässt
- keine fehlenden Werte, obwohl die Methode solche erwarten lässt
- unplausible Zeitstempel oder Erhebungsabstände
- doppelte oder fast identische Fälle
- auffällige Wiederholungen in Antwortmustern
- Variablen enthalten Werte außerhalb des zulässigen Bereichs
- beschriebene Stichprobe passt nicht zum Datensatz
- Ergebnisdarstellung ist elaborierter als die verfügbare Datenbasis
- Datenstruktur passt nicht zum beschriebenen Instrument

Prüfmethoden für Abbildungen, Diagrammen und Bildmaterial

Herkunft und Entstehung klären

Prüfen:

- Wurde die Abbildung selbst erstellt, übernommen oder KI-generiert?
- Sind KI-generierte Abbildungen als solche deutlich bezeichnet?
- Ist die Quelle angegeben?
- Sind Nutzungsrechte oder Lizenzbedingungen geklärt?
- Liegt eine Rohdatei, Originaldatei oder Zwischenversion vor?
- Sind Bearbeitungsschritte dokumentiert?
- Passt die Abbildung zur Aussage im Text?
- Ist KI-Nutzung bei generierten oder bearbeiteten Abbildungen offengelegt?

Diagramme und Datengrafiken prüfen

Prüfen:

- Stimmen Diagrammwerte mit Tabellen und Rohdaten überein?
- Sind Achsen, Einheiten und Skalen korrekt?
- Sind Gruppen, Farben und Legenden konsistent?
- Sind Fehlerbalken, Konfidenzintervalle oder Streuungsmaße erklärt?
- Werden Daten abgeschnitten, verzerrt oder irreführend dargestellt?
- Gibt es Abbildungen ohne nachvollziehbare Datenbasis?

Warnsignale:

- Diagramm sieht professionell aus, aber Werte sind nicht rückführbar
- Legende passt nicht zur Darstellung
- Achsenbeschriftungen sind ungenau oder fehlerhaft
- visuelle Aussage widerspricht den Zahlen
- Abbildungen enthalten erfundene Kategorien oder Variablen
- Datenpunkte wirken regelmäßig oder künstlich verteilt

Wissenschaftliche und medizinische Abbildungen prüfen

Prüfen:

- Ist die Bildquelle dokumentiert?
- Sind Maßstab, Modalität, Färbung, Gerät oder Untersuchungssetting angegeben?

- Passen Bild und Legende fachlich zusammen?
- Sind Wiederholungen, kopierte Bildbereiche oder unplausible Muster erkennbar?
- Sind Kontrast, Ausschnitt oder Bearbeitungsschritte offengelegt?
- Können Studierende erklären, wie die Abbildung entstanden ist?

Warnsignale:

- anatomisch, biologisch oder technisch unplausible Details
- unlesbare oder sinnlose Beschriftungen
- fehlerhafte Fachbegriffe in der Grafik
- Bilddetails widersprechen der Legende
- gleichförmige Muster, die bei realen Daten unwahrscheinlich sind
- fehlender Maßstab bei Abbildungen, bei denen er fachlich erforderlich ist

Übersicht: Merkmale KI-generierter Inhalte

Bitte beachten Sie, keines der folgenden Merkmale ist für sich genommen ein verlässlicher Nachweis für KI-generierten Text. Sie können aber bei gehäuftem Auftreten einen Hinweis auf eine KI-gestützte Erstellung geben. Moderne Sprachmodelle unterscheiden sich zudem deutlich in ihrem Schreibstil von älteren Modellen, und durch menschliche Nachbearbeitung verschwinden viele dieser Merkmale.

1. Auffällige Gedankenstriche

- **Em-Dash:** —
- **En-Dash:** –
- Konstruktionen wie:
„Das Ergebnis ist bemerkenswert — allerdings nicht überraschend.“
- Besonders auffällig kann eine sehr häufige Verwendung des langen Gedankenstrichs sein, vor allem in deutschsprachigen Texten, in denen Autorinnen und Autoren normalerweise eher Kommas, Klammern oder den normalen Bindestrich verwenden würden.

Hinweiswert: mittel bis relativ hoch, wenn der individuelle Schreibstil der Person sonst keine Gedankenstriche dieser Art enthält.

2. Typografisch „perfekte“ Anführungszeichen

Häufig erscheinen automatisch korrekte typografische Zeichen:

- “...”
- ‘...’
- „...“
- ,...‘

Auffällig kann insbesondere ein Wechsel zwischen englischen und deutschen typografischen Anführungszeichen sein. (Deutsch: „Text“; bei einem Zitat im Zitat: ,Text‘; Englisch: “Text” bei einem Zitat im Zitat: ‘Text’).

3. Überdurchschnittlich viele Doppelpunkte

KI-Texte verwenden häufig Konstruktionen wie:

- „Dabei sind drei Aspekte entscheidend:“
- „Das Ergebnis ist eindeutig:“
- „Besonders relevant ist folgender Punkt:“
- „Die zentrale Erkenntnis lautet:“

Oft folgt auf den Doppelpunkt eine Aufzählung.

4. Häufige Semikolons

;

Sprachmodelle verwenden teilweise Semikolons regelmäßiger als viele alltägliche menschliche Schreibende:

„Die Methode ist effizient; ihre klinische Relevanz bleibt jedoch unklar.“

Das ist insbesondere auffällig, wenn Semikolons im übrigen Schreibstil einer Person praktisch nicht vorkommen.

5. Sehr regelmäßige Klammerkonstruktionen

Beispiele:

- „... die Ergebnisse (insbesondere in der Subgruppe) zeigen ...“
- „... die Methode (siehe oben) ermöglicht ...“
- „... KI-Systeme (z. B. ChatGPT, Claude oder Gemini) ...“

KI tendiert dazu, Nebeninformationen sauber in Klammern einzubauen.

6. Typische Markdown-Zeichen

Ein besonders deutlicher Hinweis auf direkt aus einem Chatbot übernommenen Text können **stehen gebliebene Markdown-Zeichen** sein:

- ****Text**** – Fettdruck
- **Text** – Kursivschrift
- **#** Überschrift
- **##** Unterüberschrift
- **###** Unterüberschrift
- `Begriff`
- --- – Trennlinie

- > - Zitat
- [Text](URL) - Markdown-Link
- | Spalte 1 | Spalte 2 | - Markdown-Tabelle

Solche Artefakte sind wesentlich informativer als normale Satzzeichen.

7. Sehr systematische Bullet-Point-Strukturen

Typisch sind Aufzählungen wie:

- Erstens: ...
- Zweitens: ...
- Drittens: ...

oder

1. Ausgangslage
2. Problem
3. Lösung
4. Schlussfolgerung

Besonders auffällig ist die Kombination von **fett gesetztem Stichwort + Doppelpunkt + Erklärung**:

- **Relevanz:** ...
- **Methodik:** ...
- **Limitationen:** ...
- **Ausblick:** ...

8. Dreierstrukturen

LLMs strukturieren Argumente auffallend häufig in **Dreiergruppen**:

„Die Methode ist schnell, zuverlässig und kostengünstig.“

oder:

„Drei Aspekte sind entscheidend: Transparenz, Nachvollziehbarkeit und Verantwortung.“

Einzelne Dreierstrukturen sind völlig normal; ihre permanente Wiederholung kann jedoch charakteristisch wirken.

9. Sehr regelmäßige Zwischenüberschriften

Beispielsweise:

Hintergrund
Zentrale Herausforderungen
Mögliche Lösungsansätze
Praktische Implikationen
Fazit

KI produziert häufig eine beinahe „lehrbuchartige“ Gliederung, selbst wenn der Text eigentlich keine solche Struktur benötigt.

10. Pfeile und andere Visualisierungssymbole

Gelegentlich werden Zusammenhänge mit Zeichen dargestellt wie:

- →
- ?
- ?
- ✓
- ?
- ?
- ??
- ?

Beispielsweise:

Datenerhebung → Analyse → Interpretation → Publikation

Besonders Emojis wie ?, ?? oder ? sind für den typischen Chatbot-Stil charakteristisch, sofern sie nicht bewusst zum Kommunikationsformat gehören.

11. Ungewöhnliche Unicode-Zeichen

Beim direkten Kopieren aus KI-Systemen können Zeichen auftreten, die auf der Tastatur normalerweise nicht bewusst eingegeben werden:

- geschütztes Leerzeichen
- schmales geschütztes Leerzeichen
- typografischer Apostroph ’
- geschützter Bindestrich -
- En-Dash –

- Em-Dash —
- Ellipse ... statt drei Punkten ...

Solche Zeichen können bei einer **technischen Textanalyse** interessant sein. Sie entstehen allerdings auch durch Word, Webseiten und automatische Textformatierung.

12. Sehr konsequente Parallelkonstruktionen

Beispielsweise:

„Die erste Studie untersucht ... Die zweite Studie analysiert ... Die dritte Studie bewertet ...“

oder:

„Auf individueller Ebene ... Auf institutioneller Ebene ... Auf gesellschaftlicher Ebene ...“

Sprachmodelle produzieren besonders gerne sprachlich symmetrische Konstruktionen.

13. Häufige Kontrastkonstruktionen

Typische Formulierungen sind:

- „nicht nur ..., sondern auch ...“
- „einerseits ..., andererseits ...“
- „während ..., zeigt sich zugleich ...“
- „Obwohl ..., ist dennoch ...“
- „Dabei geht es nicht darum, ..., sondern vielmehr darum, ...“

Eine hohe Dichte solcher rhetorisch ausgewogenen Konstruktionen kann KI-typisch wirken.

14. Charakteristische Einleitungssphrasen

Häufig finden sich Formulierungen wie:

- „Es ist wichtig zu beachten, dass ...“
- „Dabei ist zu berücksichtigen, dass ...“
- „In diesem Zusammenhang ist hervorzuheben, dass ...“
- „Ein zentraler Aspekt ist ...“
- „Von besonderer Bedeutung ist ...“
- „Vor diesem Hintergrund ...“
- „Im Folgenden werden ...“

- „Grundsätzlich lässt sich festhalten ...“

Sie sind einzeln völlig normal. KI-Texte enthalten sie jedoch teilweise in ungewöhnlich hoher Frequenz.

15. Charakteristische Schlussformulierungen

Besonders häufig:

- „Zusammenfassend lässt sich sagen, dass ...“
- „Insgesamt zeigt sich, dass ...“
- „Abschließend ist festzuhalten, dass ...“
- „Letztlich wird deutlich, dass ...“
- „Dies unterstreicht die Bedeutung von ...“
- „Damit wird deutlich, dass ...“

Typisch ist auch ein letzter Satz, der die vorherigen Aussagen noch einmal **allgemein und positiv zusammenfasst**, ohne neue Information zu enthalten.

16. Inhaltlich wenig informative Metasätze

Beispiele:

- „Dieses Thema ist von großer Bedeutung.“
- „Diese Frage ist von besonderer Relevanz.“
- „Die Ergebnisse liefern wichtige Erkenntnisse.“
- „Dies eröffnet neue Perspektiven.“
- „Weitere Forschung ist erforderlich.“
- „Die Thematik ist komplex und vielschichtig.“

Solche Sätze sind nicht falsch, tragen aber häufig wenig konkrete Information. **Eine hohe Dichte solcher „glatten“, aber informationsarmen Aussagen ist wesentlich aussagekräftiger als einzelne Satzzeichen.**

17. Übermäßig ausgewogene Argumentation

KI formuliert häufig nach dem Schema:

Vorteil → Einschränkung → ausgewogene Schlussfolgerung

Beispiel:

„KI bietet erhebliche Potenziale zur Effizienzsteigerung. Gleichzeitig bestehen Herausforderungen hinsichtlich Datenschutz und Transparenz. Entscheidend ist daher ein verantwortungsvoller und reflektierter Einsatz.“

Diese fast perfekte argumentative Balance ist ein typisches LLM-Muster.

18. Ungewöhnlich gleichmäßige Satz- und Absatzstruktur

Mögliche Auffälligkeiten sind:

- fast alle Absätze ähnlich lang;
- jeder Absatz beginnt mit einer klaren Kernaussage;
- danach folgen zwei bis drei Erläuterungssätze;
- der letzte Satz fasst den Absatz zusammen oder bildet eine Überleitung;
- kaum Satzfragmente oder stilistische Brüche;
- wenige wirklich ungewöhnliche Formulierungen.

Der Text wirkt dadurch häufig **ungewöhnlich sauber, kontrolliert und gleichförmig**.

19. „Pseudo-Präzision“

Sprachmodelle produzieren gerne Formulierungen wie:

- „insbesondere“
- „maßgeblich“
- „wesentlich“
- „zentral“
- „relevant“
- „differenziert“
- „systematisch“
- „nachvollziehbar“
- „evidenzbasiert“
- „ganzheitlich“
- „robust“

Problematisch wird dies, wenn diese Begriffe häufig auftreten, ohne dass konkretisiert wird, **was** beispielsweise „robust“ oder „systematisch“ bedeutet.

20. Verdächtige Literatur- und Zitationsmuster

Bei wissenschaftlichen Texten sollte besonders auf Folgendes geachtet werden:

- auffällig gleichmäßig über Absätze verteilte Zitate
- zahlreiche Quellen hinter einer einzelnen allgemeinen Aussage; sogenannte **Reference Clusters**, z. B. [12–19] ohne das einzelne Referenzen erklärt werden
- Quellen, die die konkrete Aussage nur teilweise unterstützen
- formal plausible, aber nicht existente Literatur
- falsche DOI
- falsche Seitenzahlen
- richtige Publikation, aber falscher Inhalt
- sehr aktuelle oder sehr passende Literatur ohne erkennbare Suchstrategie.

Dies ist bei wissenschaftlichen Arbeiten häufig **aussagekräftiger als die sprachliche Oberfläche**.

Praktische Einordnung

Für die Beurteilung eines Textes würde ich die Merkmale ungefähr so gewichten:

Schwache Hinweise:

Semikolon, Doppelpunkt, typografische Anführungszeichen, Gedankenstrich, Ellipse.

Mittlere Hinweise:

ungewöhnlich viele Em-Dashes, Dreierstrukturen, stereotype Übergangsformulierungen, extrem gleichmäßige Absätze, übermäßige Zwischenüberschriften, informationsarme Metasätze.

Stärkere Hinweise:

stehen gebliebenes Markdown, typische Chatbot-Formatierungen, plötzlich gegenüber früheren Arbeiten stark veränderter Schreibstil, ungewöhnliche Unicode-Artefakte in Kombination mit weiteren Merkmalen, erfundene oder nicht passende Referenzen sowie eine Kombination vieler der oben genannten sprachlichen Muster.

Entscheidend ist daher nicht ein einzelnes „KI-Satzzeichen“, sondern das Muster. Eine seriöse Beurteilung sollte möglichst den aktuellen Text mit früheren authentischen Textproben derselben Person vergleichen und zusätzlich Inhalt, Quellen, Argumentationsführung und Entstehungsprozess berücksichtigen.