



Medienimpulse
ISSN 2307-3187
Jg. 64, Nr. 3, 2026
doi: 10.21243/mi-26-03-08
Lizenz: CC-BY-NC-ND-3.0-AT

Wenn KI Lernende beurteilt. Revidierbare Rückmeldungen zwischen Systemausgabe, Aneignung und gespeicherter Systemannahme

Eik Niederlohmman

Generative KI kann Lernenden nicht nur Texte kommentieren, sondern auch Fähigkeiten oder Lernwege zuschreiben. Wird eine solche Aussage berichtet, bleibt oft unklar, was sich geändert hat: der Wortlaut, was die lernende Person daraus ableitet oder eine gespeicherte Annahme, die das System später nutzt. Der Beitrag verbindet Forschung zu KI-Feedback, Feedbackkompetenz, offenen Lernendenmodellen und Anfechtbarkeit. Er schlägt eine nachvollziehbare Korrekturspur und Revidierbarkeit als noch zu prüfendes Qualitätskriterium vor. Ein konstruiertes Unterrichtsszenario zeigt vier gleichwertige Möglichkeiten: bestäti-

gen, differenzieren, zurückweisen oder vertagen. Lehrkräfte schaffen fachliche Prüfgelegenheiten; Institutionen und Anbieter sichern technische und beschwerdebezogene Wege. So wird epistemische Würde praktisch: Lernende können begründet widersprechen, ohne sich persönlich offenlegen oder Nachteile befürchten zu müssen.

Generative AI can do more than comment on learners' work: it can also make claims about their abilities or possible educational pathways. When such a claim is corrected, it may remain unclear what has changed—the wording, what the learner takes from it, or a stored assumption that the system later reuses. This problem-oriented conceptual synthesis connects research on AI feedback, feedback literacy, open learner models, and contestability. It proposes a traceable correction pathway and revisability as a quality criterion that still requires empirical testing. A constructed classroom scenario translates the idea into four equally permissible responses: confirm, differentiate, reject, or defer. Teachers create opportunities for subject-based review; institutions and providers secure technical and complaint routes. In this way, epistemic dignity becomes practical: learners can disagree with reasons without having to disclose personal meaning or fear disadvantage.

1. Wenn KI etwas über Lernende aussagt

1.1 Eine alltägliche, aber folgenreiche Frage

Eine Schülerin lässt drei argumentative Texte durch ein dialogisches KI-System kommentieren. Die Ausgabe lautet: „Aus diesen Texten schließe ich, dass deine Stärke eher im persönlichen Ausdruck als in analytischer Argumentation liegt“. Die Klasse prüft die Texte mit einem transparenten Raster. Aus dem globalen Urteil

wird eine engere Aussage: „In zwei der drei Texte bleiben zentrale Behauptungen noch ohne Textbeleg“. Damit ist der Wortlaut berichtigt. Offen bleibt, was die Schülerin daraus für ihr weiteres Lernen ableitet und ob ein System mit Gedächtnis die frühere Annahme später erneut verwendet. Die Szene ist konstruiert. Sie macht ein Gestaltungsproblem sichtbar, ohne dessen Häufigkeit oder psychische Wirkung zu behaupten.

Mediendidaktik fragt nicht nur nach einem Werkzeug. Sie gestaltet Aufgaben, Materialien, Rollen, Interaktionen und institutionelle Bedingungen so, dass ein Bildungsproblem bearbeitbar wird (Kerres, 2024). Generative KI verändert dabei nicht jede didaktische Frage von Grund auf; sie verschiebt aber Anforderungen an Kompetenzen, Prüfungen und Lehr-Lern-Methoden (Klar & Schleiss, 2024). Wenn KI-Rückmeldungen Aussagen über Lernende enthalten, gehören deshalb auch ihre Bedeutung im Lernprozess und erreichbare Korrekturwege zum Arrangement.

KI-gestützte Rückmeldung kann neben Aufgaben- und Prozessebene auch die lernende Person bewerten (Ba et al., 2025). Rückmeldung kann fachbezogene Selbstkonzepte berühren; die einschlägige schulische Evidenz stammt jedoch nicht aus generativen KI-Settings (Van der Beek et al., 2025). Entscheidend ist daher zunächst eine nüchterne Frage: Was genau soll revidierbar sein?

1.2 Beitrag und Lernziel

Arbeiten zu epistemischem Auslagern, digitaler Selbstbestimmung, algorithmischer Schließung und pädagogischer Souveräni-

tät zeigen bereits, wie KI Urteilspraktiken und Medienbildung mitprägt (Filk, 2025, 2026; Swertz & Knaus, 2026; Winder, 2026). Feedbackkompetenz betrifft die Nutzung von Rückmeldung, Anfechtbarkeit Eingriffe in Systeme oder Entscheidungen. Offene Lernendenmodelle machen Modellannahmen sichtbar und teils bearbeitbar.

Der Zusatznutzen liegt in einer gemeinsamen Korrekturspur. Ein Fall gilt als nachvollziehbar bearbeitet, wenn erkennbar ist, welcher Gegenstand geändert wurde, ob eine spätere Wiederverwendung geprüft werden muss und wer für einen offenen nächsten Schritt zuständig bleibt. Korrekturdurchgängigkeit meint also keine automatische Übernahme. Lernende müssen eine geänderte Systemaussage nicht annehmen. Die Spur zeigt vielmehr: Was wurde geändert, was blieb unberührt und wer kann weiterhelfen?

Als Lernziel steht eine beobachtbare Praxis im Vordergrund. Lernende können eine Aussage über sich erkennen, die angeführten Texte und Kriterien prüfen, eine andere Erklärung erwägen und begründet entscheiden, ob sie die Aussage bestätigen, differenzieren, zurückweisen oder vorläufig offenlassen.

2. Wie die Literatur ausgewählt und verbunden wurde

2.1 Problemorientierte Synthese

Der Beitrag entwickelt die Unterscheidung durch eine problemorientierte konzeptionelle Synthese. Berücksichtigt wurden deutsch- und englischsprachige Arbeiten bis 22. Juli 2026 aus dem MEDI-

ENIMPULSE-Archiv, Verlags- und Repositoriumssuchen, offiziellen Leitlinien sowie Rückwärts- und Vorwärtssuchen. Ältere Grundlagenarbeiten ergänzen Feedback, Selbstkonzept, Attribution, offene Lernendenmodelle, Anfechtbarkeit und Reviewmethodik.

Die Quellen wurden danach geordnet, ob sie Systemausgaben, das Deuten von Lernenden oder gespeicherte Annahmen beleuchten. Der Text unterscheidet Befunde, theoretische Positionen, konzeptionelle Brücken, normative Setzungen und Praxisvorschläge (Jaakkola, 2020; Snyder, 2019).

2.2 Reichweite

Es handelt sich nicht um ein systematisches Review. Datenbankexport, unabhängiges Doppelscreening und PRISMA-Selektion erfolgten nicht. Vollständigkeits- und Wirksamkeitsansprüche wären deshalb unangebracht; die Darstellung orientiert sich an Qualitätskriterien für narrative Übersichten (Baethge et al., 2019). Revidierbarkeit wird als prüfbarer Vorschlag entwickelt, nicht als bereits validierter Standard.

3. Was genau revidiert werden soll

3.1 Drei verschiedene Gegenstände

Als lernendenbezogene KI-Aussage gilt hier eine Systemausgabe zu Fähigkeiten, Schwierigkeiten, Strategien, Fortschritt oder Bildungsoptionen. „Dieser Text hat wenige Übergänge“ besitzt einen anderen Status als „Du bist nicht analytisch“ oder „Ein technischer Bildungsweg passt nicht zu dir“. Aus Sicht dieses Beitrags sollten

Belege und menschliche Prüfung umso stärker sein, je globaler, stabiler oder entscheidungsnäher eine Aussage formuliert ist. Attributionstheorie hilft dabei zu fragen, ob Schwierigkeiten als veränderbar und kontrollierbar oder als feste Eigenschaft erscheinen (Weiner, 1985).

Aneignung bezeichnet die Bedeutung, die Lernende einer Aussage für ihr weiteres Lernen geben. Externe Herkunft ist nicht automatisch Fremdbestimmung; Lernen und Handlungsfähigkeit sind sozial und technisch vermittelt (Mhlambi & Tiribelli, 2023). Ob Lernende eine Aussage neu bewerten, entscheiden sie selbst. Lehrkräfte können Gründe und Gegenbelege zugänglich machen, aber keine „richtige“ Selbstdeutung verordnen.

Eine gespeicherte Systemannahme liegt nur vor, wenn ein System Informationen über eine Person speichert und später zur Personalisierung nutzt. Flüchtige Ausgabe, persönliche Deutung und technisch gespeicherte Annahme sind daher nicht dasselbe.

Epistemische Handlungsfähigkeit meint, Wissen zu beurteilen, hervorzubringen, anzuwenden und zu verändern (Nieminen & Ketonen, 2024). Feedback kann dabei als Wissenspraxis verstanden werden, in der Lernende als legitime Wissende angesprochen werden (Nieminen et al., 2026). Unter epistemischer Würde wird hier – als normative Setzung, nicht als Wirkmechanismus – der Anspruch verstanden, als urteilsfähige Person behandelt zu werden: Lernende können eine sie betreffende Zuschreibung bestreiten, nach Gründen fragen und eine Berichtigung verlangen, ohne

dadurch schlechter gestellt zu werden (in Anlehnung an Abimbo-la, 2025).

3.2 Anschluss an benachbarte Konzepte

Tabelle 1 ordnet den Vorschlag ein. Die Konzepte stehen nicht in Konkurrenz; sie beantworten verschiedene Teile derselben Gestaltungsaufgabe:

<i>Konzept</i>	<i>Was es leistet</i>	<i>Leitfrage</i>	<i>Verbleibende Frage im vorliegenden Zusammenhang</i>
Kritische KI-Kompetenz	Macht Daten, Systemgrenzen und Macht zum Lerngegenstand.	Wie lässt sich KI verstehen, beurteilen und mitgestalten?	Wie bleibt eine Berichtigung über Ausgabe, Deutung und Speicherung nachvollziehbar?
Feedbackkompetenz	Beschreibt, wie Lernende Rückmeldung verstehen, bewerten und nutzen.	Was lässt sich daraus für weiteres Lernen gewinnen?	Wie greifen pädagogische, technische und institutionelle Wege ineinander?
Anfechtbarkeit und Abhilfe	Ermöglicht, Ergebnisse oder Folgen zu bestreiten und zu ändern.	Wer kann mit welchen Gründen eingreifen?	Wie wird zugleich sichtbar, was Lernende daraus ableiten?
Offene Lernendenmodelle	Machen Modellannahmen einsehbar und teils aushandelbar oder editierbar.	Wie können Lernende ein Modell ihrer Kenntnisse prüfen und mitbearbeiten?	Sie setzen ein explizites Modell voraus und verfolgen selten die ganze Spur einer generativen Aus-

			gabe.
Pädagogische Souveränität	Schützt Subjektstatus und professionelle Urteilskraft.	Was darf nicht an KI abgegeben werden?	Welche Stelle braucht im Einzelfall einen Korrekturweg?
Nachvollziehbare Korrekturspur	Verbindet Ausgabe, Aneignung und spätere Wiederverwendung.	Was wurde geändert, was blieb offen und wer ist zuständig?	Ihr praktischer Zusatznutzen muss empirisch geprüft werden.

Tabelle 1: Benachbarte Konzepte und die verbleibende Korrekturfrage

Anmerkung. Eigene Darstellung. Die letzte Zeile bezeichnet keinen neuen Kompetenzrahmen, sondern verbindet bereits bearbeitete Korrekturfragen (Alfrink et al., 2023; Carless & Boud, 2018; Hooshyar et al., 2020; Karimi et al., 2022; Robles Mucho et al., 2026; Veldhuis et al., 2025 [CC-BY-SA])

4. Was die Evidenz zeigt – und was offen bleibt

4.1 Rückmeldung, Selbstkonzept und KI-Rezeption

Eine systematische Übersicht unterscheidet KI-Feedback auf Aufgaben-, Prozess-, Selbstregulations- und Selbstebene (Ba et al., 2025). Eine schulische Experimentalstudie zeigt, dass Mathematikfeedback Emotionen, Leistung und Selbstkonzept abhängig von Vorleistung und bestehendem Selbstkonzept beeinflussen kann; generative KI wurde dort nicht untersucht (Van der Beek et al., 2025). Akademisches Selbstkonzept und Leistung wirken wechselseitig aufeinander (Marsh & Martin, 2011). Rückmeldung hilft vor allem dann beim Lernen, wenn sie zu Aufgabe, Prozess und nächstem Schritt passt (Hattie & Timperley, 2007).

In einer Hochschulfbefragung wurde generatives KI-Feedback als schnell, zugänglich und sozial risikoärmer erlebt; Lehrkraftfeed-

back galt insgesamt als hilfreicher und vertrauenswürdiger (Henderson et al., 2025). In nichtpädagogischen Frage-Antwort-Aufgaben überschätzten Erwachsene teils die Genauigkeit von LLM-Antworten (Steyvers et al., 2025). Foidl (2026) beschreibt sprachliche Belebtheitskonstruktionen. Watson und Ennion (2026) argumentieren, generative KI könne Bedeutungen von Anstrengung, Fehlern und Kompetenz mitrahmen. Das begründet Prüfung, nicht die Annahme besonderer psychischer Macht.

4.2 Grenzen und konkurrierende Erklärungen

Direkte Längsschnittbefunde dazu, dass generative KI fachbezogene Selbstdeutungen oder Bildungsentscheidungen dauerhaft verändert, fehlen in der ausgewerteten Literatur. Ebenso ist nicht geklärt, ob KI-Rückmeldung stärker wirkt als inhaltlich und sozial vergleichbares Feedback durch Lehrkräfte oder Peers. Vorwissen, bestehendes Selbstkonzept, Noten, das Framing der Lehrkraft, Peers und eigene Leistungserfahrungen können wichtiger sein als die technische Quelle.

KI-spezifisch sind daher weniger pauschale Zuschreibungen als solche. Spezifisch können Skalierung, Dialoggeschwindigkeit, Personalisierung, begrenzte Nachvollziehbarkeit und – bei bestimmten Systemen – Speicherung sein. Sie werden im Folgenden als Bedingungen für Prüf- und Korrekturwege behandelt, nicht als nachgewiesene Wirkungskette.

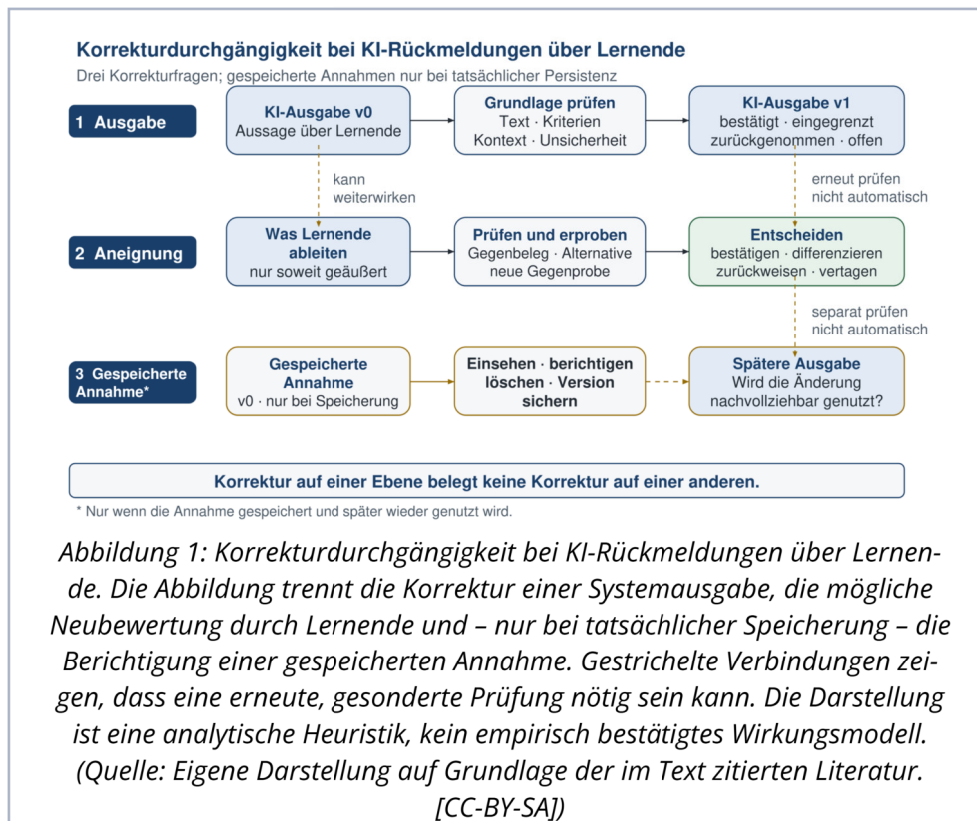
5. Eine nachvollziehbare Korrekturspur

5.1 Ausgabe, Aneignung und gespeicherte Annahme

Die erste Frage betrifft eine konkrete Systemausgabe. Sie wird überprüfbar, wenn erkennbar ist, welche Lernprodukte, Kriterien, Vergleichsmaßstäbe und Kontexte die Aussage stützen, welche Unsicherheit besteht und welche Gegenbelege zu einer Änderung führen. Eine neu formulierte Ausgabe zeigt noch nicht, was Lernende aus ihr ableiten.

Die zweite Frage betrifft eine mögliche Neubewertung durch die lernende Person. Feedback wird erst durch Verstehen, Bewerten und weitere Nutzung Teil des Lernens (Carless & Boud, 2018). Schweigen, erneute Nutzung oder einzelne Eingaben erlauben keinen sicheren Schluss auf persönliche Bedeutung. Eine Revision bleibt eine Möglichkeit, keine pädagogische Erwartung.

Die dritte Frage betrifft gespeicherte Annahmen, die spätere Ausgaben beeinflussen. Offene Lernendenmodelle reichen von einsehbaren bis zu editierbaren oder aushandelbaren Varianten und können Reflexion, Selbstregulation und Berichtigung unterstützen (Hooshyar et al., 2020; Robles Mucho et al., 2026). Im Umfeld generativer KI wird neu verhandelt, wie solche Modelle transparent, dialogisch und kontrollierbar bleiben (Chounta et al., 2026). Die Korrekturspur beginnt jedoch bereits bei einer einzelnen generativen Ausgabe und fragt zusätzlich, ob eine bestätigte Berichtigung in späterer Personalisierung ankommt:



5.2 Wann die dritte Frage überhaupt entsteht

Bei Einzelabfragen ohne Speicherung sind vor allem Ausgabe und Aneignung relevant. In einem Dialog kann eine frühere Formulierung den weiteren Austausch prägen, ohne dauerhaft gespeichert zu werden. Erst Kontogedächtnis, internes Nutzerprofil oder adaptives Lernendenmodell eröffnen die dritte Frage. Ein offenes Lernendenmodell ist dabei nicht mit einem verborgenen Plattformprofil gleichzusetzen.

Das angemessene Ergebnis muss keine Änderung sein. Eine Aussage kann bestätigt werden, wenn tragfähige Gründe und weitere

Aufgaben sie stützen. Gleichberechtigt bedeutet, dass keine Option erzwungen wird. Welche Option fachlich trägt, hängt dennoch von Gründen und Belegen ab.

Besonders wichtig ist die Korrekturspur bei globalen, prognostischen, wiederholten oder selektionsnahen Aussagen. Diskriminierende Zuschreibungen, verpflichtende Nutzung und folgenreiche Entscheidungen brauchen darüber hinaus menschliche Prüfung, Datenschutz und erreichbare Beschwerdewege. Eine Reflexionsaufgabe allein reicht dann nicht.

6. Wie Revidierbarkeit im Lernarrangement praktisch wird

6.1 Mindestanforderungen und Zuständigkeit

Als mediendidaktisches Kriterium bezieht sich Revidierbarkeit auf das gesamte Lernarrangement: Aufgaben, Materialien, Rollen, Interaktionswege und institutionelle Bedingungen sollen einen begründeten Widerspruch tragen können. Tabelle 2 benennt dafür beobachtbare Mindestanforderungen. Verantwortung wird nach tatsächlicher Kontrollmacht verteilt. Lernende müssen kein opakes System reparieren:

<i>Prüfbereich</i>	<i>Woran der Weg erkennbar ist</i>	<i>Warnsignal</i>	<i>Zuständigkeit</i>	<i>Schutzanforderung</i>
Systemausgabe	Lernprodukt, Kriterien, Kontext und Unsicherheit sind lokal er-	Globales Personenurteil ohne prüfbare Grundlage oder Korrek-	Anbieter; Institution bei Auswahl und Konfiguration.	Verständliche Darstellung; keine Scheingenauigkeit.

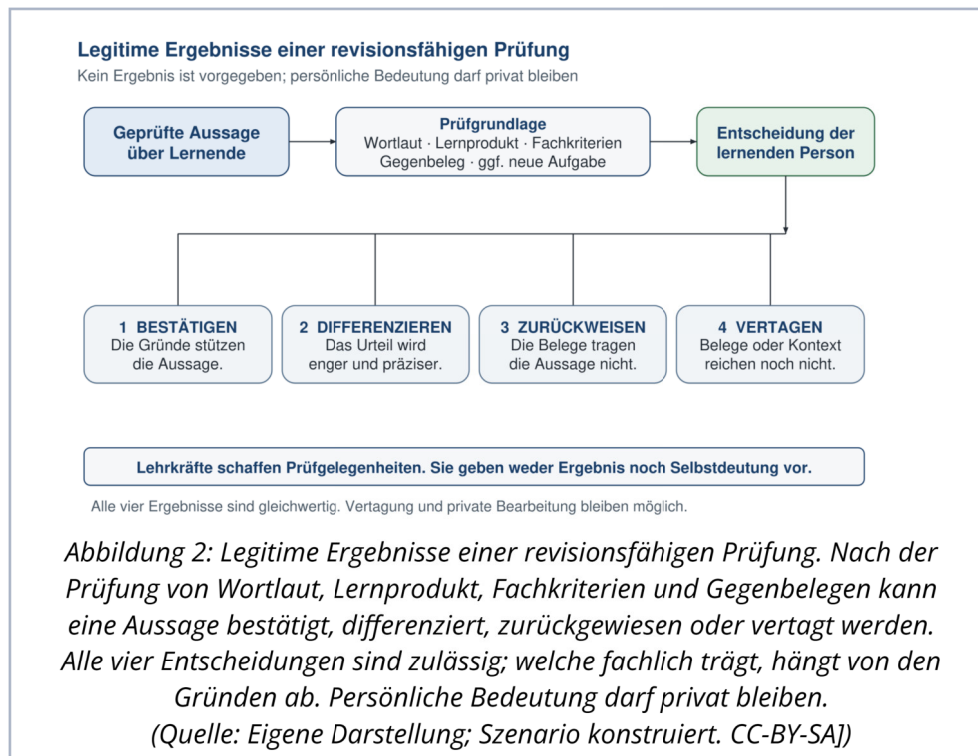
	kennbar.	turoption.		
Lernaufgabe	Gegenbeleg oder andere Erklärung kann anhand fachlicher Kriterien geprüft werden.	Nur erneutes Prompten; keine fachliche Gegenprobe.	Lehrkraft und Institution.	Kein Zwang zur persönlichen Selbstoffenlegung.
Entscheidung	Bestätigen, differenzieren, zurückweisen oder vertagen bleiben zulässig.	Eine bestimmte Selbstdeutung wird erwartet oder benotet.	Lernende Person; Lehrkraft und Institution schützen Freiwilligkeit.	Kein Nachteil durch Widerspruch oder Vertagung.
Gespeicherte Annahme*	Die Annahme ist einsehbar, kann berichtigt oder gelöscht werden, und ihre spätere Nutzung lässt sich prüfen.	Unklare Speicherung oder fortgesetzte Nutzung der alten Annahme.	Anbieter und Institution.	Datenschutz, Barrierefreiheit, Nichtdiskriminierung.
Institutioneller Weg	Zuständigkeit, menschliche Prüfung, Frist und Rückmeldung sind benannt.	Niemand ist erreichbar; Folgen bleiben ungeklärt.	Schule, Träger und Anbieter.	Korrektur ohne Teilhabe- oder Bewertungsnachteil.

Tabelle 2: Mindestanforderungen für revidierbare KI-Rückmeldungen über Lernende (CC-BY-SA)

7.1 Vier legitime Ergebnisse

Im konstruierten Deutschunterricht entscheidet die Schülerin zunächst, ob sie die globale Aussage überhaupt prüfen möchte. Grundlage sind der Wortlaut, ihre drei Texte und ein transparentes Raster. Die Lehrkraft fragt nicht, welche Selbstbeschreibung „richtig“ ist, sondern welche Textstellen die engere Aussage stützen oder widerlegen. Eine kurze Prüfung kann zu vier legitimen Ergebnissen führen: Die Aussage wird bestätigt, differenziert, zurückgewiesen oder bis zu weiteren Belegen vertagt.

Eine neue Schreibaufgabe oder zusätzliche Rückmeldung kann helfen, ist aber keine Pflichtsequenz. Auch die Lehrkraft gibt das Ergebnis nicht vor. Sie macht den Unterschied zwischen Textbefund, globaler Fähigkeitszuschreibung und persönlicher Bedeutung sichtbar:



7.2 Beobachtbare Praxis statt Haltungsprüfung

Forschung zum KI-Einsatz im Unterricht ist umfangreicher als Forschung zur systematischen KI-bezogenen Lehrkräftebildung; schulische KI-Kompetenz bleibt begrifflich und methodisch heterogen (Feng & Carolus, 2026; Tan et al., 2025). Die folgenden Funktionen bilden daher keine neue Kompetenzskala. Sie übersetzen vorhandene Anforderungen auf diesen Fall: Aussagen über Lernende erkennen; Aussage, Beleg und Zuschreibung trennen; eine proportionale fachliche Gegenprobe gestalten; bei Profil-, Datenschutz- oder Diskriminierungsfragen die zuständige Stelle einschalten.

In der Lehrkräftebildung lässt sich das als kurze Mikroübung erproben. Eine Person liest einen KI-Output. Eine zweite markiert in 90 Sekunden Textbefund, Personenurteil und fehlende Begründung. Danach formuliert sie eine fachliche Rückfrage und benennt, wer für den nächsten Schritt zuständig ist. Die Rückmeldung bezieht sich nur auf diese beobachtbaren Schritte, nicht auf Haltung oder Persönlichkeit. Erst danach steigt die Schwierigkeit: von lokalem Textfeedback zu gespeicherten Profilannahmen oder selektionsnahen Empfehlungen. Diese Mikroübung ist ein Vorschlag für die Praxis, kein erprobtes Curriculum.

8. Grenzen und Forschungsagenda

8.1 Evidenz, Ungleichheit und Macht

Der Beitrag zeigt nicht, wie häufig Korrekturen zwischen den drei Gegenständen abbrechen. Viele Studien zu generativem KI-Feedback stammen aus Hochschulen; ihre Übertragbarkeit auf Schule ist begrenzt. Es ist auch nicht belegt, dass Jugendliche generell stärker beeinflussbar sind. Prüf- und Beschwerdewege kosten Zeit, Sprache, Aufmerksamkeit und Vertrauen. Wenn diese Ressourcen ungleich verteilt sind, kann ein gut gemeintes Revisionsangebot bestehende Unterschiede verstärken.

Die Macht ist asymmetrisch verteilt. Anbieter kontrollieren Systemdesign und technische Speicherung, Institutionen Beschaffung und Verfahren, Lehrkräfte das konkrete Lernarrangement. Lernende kontrollieren diese Ebenen nicht im gleichen Maß. Epis-

temische Vertrauenswürdigkeit ist deshalb eine geteilte, aber ungleich verteilte Verantwortung (Tanchuk & Taylor, 2025). Eine diskriminierende Zuschreibung darf nicht zur privaten Reflexionsaufgabe der betroffenen Person werden; zu korrigieren sind zuerst System, Bewertungspraktik oder institutionelle Nutzung.

Plattform- und Geschäftsmodelle können Personalisierung, Speicherung und Anbieterwechsel vorstrukturieren (Schmidt, 2025; Tanchuk, 2026). Der Beitrag unterstellt keine Manipulationsabsicht und bewertet kein einzelnes Produkt. Er fragt konkret: Was wird gespeichert, wer kann es sehen, ändern oder löschen, und was bewirkt eine Berichtigung in späteren Ausgaben?

8.2 Prüffragen und Revisionsbedingungen

Vier Fragen können empirische Arbeiten leiten. *Erstens*: Lassen sich eine Änderung des Wortlauts und eine Änderung dessen, was Lernende daraus ableiten, in konkreten Situationen getrennt beobachten? *Zweitens*: Verwendet ein System nach einer bestätigten Berichtigung die frühere Annahme in späteren Ausgaben weiter? *Drittens*: Unterstützen sichtbare Kriterien, Gegenbelege und alternative Erklärungen eine angemessenere Einschätzung, ohne Lernende zusätzlich zu belasten? *Viertens*: Wer trägt die zeitlichen, sprachlichen und institutionellen Kosten eines Widerspruchs?

Dafür eignen sich kontrollierte Vergleiche von KI-, Lehrkraft- und Peerfeedback, Längsschnittstudien zu fachbezogenen Selbstdeutungen, technische Prüfungen vor und nach einer Profilberichtigung sowie partizipative, barriere- und diskriminierungssensible

Nutzungsstudien. Wenn die Unterscheidung keinen praktischen Zusatznutzen besitzt oder Orientierung, Gleichheit und Freiwilligkeit verschlechtert, muss sie eingegrenzt oder aufgegeben werden.

9. Schluss: Korrektur als Qualität des Lernarrangements

Eine berichtete Systemausgabe, das, was Lernende daraus für sich ableiten, und eine aktualisierte gespeicherte Annahme sind verschiedene Vorgänge. Gute Lernarrangements machen sichtbar, welcher davon betroffen ist, welche Frage offen bleibt und wer helfen kann.

Lernende sollen Aussagen über sich bestätigen, differenzieren, zurückweisen oder vertagen können. Lehrkräfte schaffen fachliche Prüfgelegenheiten und schützen Freiwilligkeit. Institutionen und Anbieter sichern verständliche, barrierefreie und nachteilsfreie Korrekturwege. So unterstützt Revidierbarkeit eine demokratisch orientierte Medienbildung: nicht durch ein Wirkungsversprechen, sondern durch die gelebte Möglichkeit, begründet zu widersprechen und eine begründete Meinungsänderung zuzulassen.

Erklärungen

Forschungsdaten und Ethik: Der Beitrag ist eine konzeptionelle Synthese. Es wurden keine eigenen Daten erhoben; das Unterrichtsszenario ist konstruiert.

Förderung: Der Autor erhielt keine finanzielle Unterstützung für Recherche, Autorenschaft oder Veröffentlichung dieses Beitrags.

Interessenkonflikte: Der Autor erklärt, dass keine Interessenkonflikte im Zusammenhang mit diesem Beitrag bestehen.

Beitrag des Autors: EN konzipierte den Beitrag, führte die Literatursynthese durch, entwickelte Heuristik und Abbildungen und verfasste und überarbeitete das Manuskript.

Einsatz generativer KI: ChatGPT (OpenAI; Juli 2026) unterstützte Rechercheplanung und Sprachprüfung. Der Autor prüfte die Quellen im Original, traf alle inhaltlichen Entscheidungen und trägt die Verantwortung für den Beitrag. Personenbezogene oder vertrauliche Daten wurden nicht verarbeitet.

Rechte: Tabellen und Abbildungen sind eigene Darstellungen; Fremdbilder oder lizenzpflichtige Symbole wurden nicht verwendet.

Literatur

Abimbola, S. (2025). Epistemic dignity. *The Lancet*, 406(10516), 2210–2211. [https://doi.org/10.1016/S0140-6736\(25\)02212-3](https://doi.org/10.1016/S0140-6736(25)02212-3)

Alfrink, K., Keller, I., Kortuem, G., & Doorn, N. (2023). Contestable AI by Design: Towards a Framework. *Minds and Machines*, 33, 613–639. <https://link.springer.com/article/10.1007/s11023-022-09611-z>

Ba, S., Yang, L., Yan, Z., Looi, C. K., & Gašević, D. (2025). Unraveling the mechanisms and effectiveness of AI-assisted feedback in education: A systematic literature review. *Computers and Education Open*, 9, Article 100284. <https://doi.org/10.1016/j.caeo.2025.100284>

Baethge, C., Goldbeck-Wood, S., & Mertens, S. (2019). SANRA—a scale for the quality assessment of narrative review articles. *Research Integrity and Peer Review*, 4, Article 5. <https://doi.org/10.1186/s41073-019-0064-8>

Carless, D., & Boud, D. (2018). The development of student feedback literacy: Enabling uptake of feedback. *Assessment & Evaluation in Higher Education*, 43(8), 1315–1325. <https://doi.org/10.1080/02602938.2018.1463354>

Chounta, I.-A., Sheng, K., Abdelmagied, M., Nagashima, T., & Zapata-Rivera, D. (2026). Open Learner Models in the Age of Generative AI. In E. G. Blanchard, G. Chen, M. Chi & S. Isotani (Hrsg.), *Artificial intelligence in education: Late breaking results, WideAIED, practitioners, industry and policies* (Communications in Computer and Information Science, Bd. 3033, S. 61–65). Springer. https://doi.org/10.1007/978-3-032-29794-5_11

Feng, S., & Carolus, A. (2026). Artificial intelligence literacy at school: A systematic review with a focus on psychological foundations. *Computers and Education: Artificial Intelligence*, 10, Article 100551. <https://doi.org/10.1016/j.caeai.2026.100551>

Filk, C. (2025). Strengthening Digital Self-Determination. Integrating Media Ethics and Artificial Intelligence into Teacher Training for Everyday School Life. *Medienimpulse*, 63(1), 48 Seiten. <https://doi.org/10.21243/mi-01-25-29>

Filk, C. (2026). Inexistent: Algorithmische Medienkulturen und die strukturelle Schließung der Vollzugsbedingungen von Medienbildung. *Medienimpulse*, 64(2), 32 Seiten. <https://doi.org/10.21243/mi-02-26-24>

Foidl, R. (2026). Eine verschwimmende Grenze: Sprachliche Belebtheitskonstruktion in der Mensch-KI-Interaktion. *Medienimpulse*, 64(2), 38 Seiten. <https://doi.org/10.21243/mi-02-26-15>

Hattie, J., & Timperley, H. (2007). The Power of Feedback. *Review of Educational Research*, 77(1), 81–112. <https://doi.org/10.3102/003465430298487>

Henderson, M., Bearman, M., Chung, J., Fawns, T., Buckingham Shum, S., Matthews, K. E., & de Mello Heredia, J. (2025). Comparing Generative AI and teacher feedback: student perceptions of usefulness and trustworthiness. *Assessment & Evaluation in Higher Education*. Advance online publication. <https://doi.org/10.1080/02602938.2025.2502582>

Hooshyar, D., Pedaste, M., Saks, K., Leijen, Ä., Bardone, E., & Wang, M. (2020). Open learner models in supporting self-regulated learning in higher education: A systematic literature review. *Computers & Education*, 154, Article 103878. <https://doi.org/10.1016/j.compedu.2020.103878>

Jaakkola, E. (2020). Designing conceptual articles: Four approaches. *AMS Review*, 10(1–2), 18–26. <https://doi.org/10.1007/s13162-020-00161-0>

Karimi, A.-H., Barthe, G., Schölkopf, B., & Valera, I. (2022). A Survey of Algorithmic Recourse: Contrastive Explanations and Consequential Recommendations. *ACM Computing Surveys*, 55(5), Article 95. <https://doi.org/10.1145/3527848>

Kerres, M. (2024). *Mediendidaktik: Lernen in der digitalen Welt* (6. Aufl.). De Gruyter Oldenbourg. <https://doi.org/10.1515/9783111201078>

Klar, M., & Schleiss, J. (2024). Künstliche Intelligenz im Kontext von Kompetenzen, Prüfungen und Lehr-Lern-Methoden: Alte und neue Gestaltungsfragen. *MedienPädagogik: Zeitschrift für Theorie und Praxis der Medienbildung*, 58, 41–57. <https://doi.org/10.21240/mpaed/58/2024.03.24.X>

Marsh, H. W., & Martin, A. J. (2011). Academic self-concept and academic achievement: Relations and causal ordering. *British Journal of Educational Psychology*, 81(1), 59–77. <https://doi.org/10.1348/000709910X503501>

Mhlambi, S., & Tiribelli, S. (2023). Decolonizing AI ethics: Relational autonomy as a means to counter AI harms. *Topoi*, 42(3), 867–880. <https://doi.org/10.1007/s11245-022-09874-2>

Miao, F., & Holmes, W. (2023). *Guidance for generative AI in education and research*. UNESCO. <https://doi.org/10.54675/EWZM9535>

Nieminen, J. H., & Ketonen, L. (2024). Epistemic agency: A link between assessment, knowledge and society. *Higher Education*, 88(2), 777–794. <https://doi.org/10.1007/s10734-023-01142-5>

Nieminen, J. H., Walton, J., & Cobb, P. J. (2026). Feedback as an epistemic practice. *Assessment & Evaluation in Higher Education*. Advance online publication. <https://doi.org/10.1080/02602938.2026.2624483>

Robles Mucho, J. L., Andrade-Girón, D. C., Benites-Tirado, V. R., Quispe-Maquera, N. B., Ayala-Jara, C., Tito Chura, H. E., Luna-Victoria, F. M., Laura-De La Cruz, K. M., Rivera-Lozada, O., Cerna Salcedo, A. A., Jaramillo Arica, P. S., & Barboza, J. J. (2026). Open learner models and pedagogical strategies in higher education: A meta-synthesis of approaches to self-regulated learning. *Frontiers in Education*, 10, Article 1760183. <https://doi.org/10.3389/feduc.2025.1760183>

Schmidt, C. (2025). Ökonomisierung reloaded: Oder: Wie Ed-Tech Unterricht verändert. *Medienimpulse*, 63(4), 34 Seiten. <https://doi.org/10.21243/mi-04-25-10>

Snyder, H. (2019). Literature review as a research methodology: An overview and guidelines. *Journal of Business Research*, 104, 333–339. <https://doi.org/10.1016/j.jbusres.2019.07.039>

Steyvers, M., Tejeda, H., Kumar, A., Belem, C., Karny, S., Hu, X., Mayer, L. W., & Smyth, P. (2025). What large language models know and what people think they know. *Nature Machine Intelligence*, 7(2), 221–231. <https://doi.org/10.1038/s42256-024-00976-7>

Swertz, C., & Knaus, T. (2026). Keine Bildung ohne Subjekt? Ein Gespräch über KI, Lernen als Erkennen und pädagogische Souveränität. *Medienimpulse*, 64(2), 50 Seiten. <https://doi.org/10.21243/mi-02-26-07>

Tan, X., Cheng, K. S., & Ling, M. H. A. (2025). Artificial intelligence in teaching and teacher professional development: A systematic review. *Computers and Education: Artificial Intelligence*, 8, Article 100355. <https://doi.org/10.1016/j.caeai.2024.100355>

Tanchuk, N. J. (2026). AI personalized Learning and the Risk of Epistemic Consolidation. *Studies in Philosophy and Education*. Advance online publication. <https://doi.org/10.1007/s11217-026-10081-4>

Tanchuk, N. J., & Taylor, R. M. (2025). Personalized Learning with AI Tutors: Assessing and Advancing Epistemic Trustworthiness. *Educational Theory*, 75(2), 327–353. <https://doi.org/10.1111/edth.70009>

Van der Beek, J. P. J., Van der Ven, S. H. G., Van de Weijer-Bergsma, E., & Kroesbergen, E. H. (2025). How feedback affects emotions, performance and self-concept in mathematics. *Learning and Instruction*, 99, Article 102169. <https://doi.org/10.1016/j.learninstruc.2025.102169>

Veldhuis, A., Lo, P. Y., Kenny, S., & Antle, A. N. (2025). Critical Artificial Intelligence literacy: A scoping review and framework synthesis. *International Journal of Child-Computer Interaction*, 43, Article 100708. <https://doi.org/10.1016/j.ijcci.2024.100708>

Watson, S., & Ennion, M. (2026). Programmable support: Generative AI, behavioural influence, and the mediation of meaning in education. *Learning, Media and Technology*. Advance online publication. <https://doi.org/10.1080/17439884.2026.2642199>

Weiner, B. (1985). An attributional theory of achievement motivation and emotion. *Psychological Review*, 92(4), 548–573. <https://doi.org/10.1037/0033-295X.92.4.548>

Winder, G. (2026). *Wer entscheidet, was wir denken? Epistemisches Auslagern, algorithmische Infrastrukturen und die demokratische Aufgabe der Medienbildung*. *Medienimpulse*, 64(2), 26 Seiten. <https://doi.org/10.21243/mi-02-26-12>