

Inhaltsverzeichnis

**Vorwort: Wenn KI auf Fach und Sprache trifft –
Zur Einführung in den Band 9**

**FACHKOMMUNIKATION UND KI –
MODELLE, KONZEPTE UND TERMINOLOGIE**

**THORSTEN ROELCKE
Fachliche Kommunikation mit Künstlicher Intelligenz 15**

**MARGRET MUNDORF
Rechtskommunikation im Umbruch: Ein Framework
für *AI Legal Literacy* und fachkommunikative Kompetenzen
im Recht 29**

**TANJA WISSIK
Auswirkungen von KI auf Terminologearbeit
im institutionellen Kontext: Workflows, Rollen und Kompetenzen 51**

**KAROLINA SUCHOWOLEC / CHRISTIAN LANG
Terminologie und Retrieval-Augmented Generation (TermRAG) 75**

FACHKOMMUNIKATIONSFORSCHUNG MIT UND ZU KI

WALDEMAR NAZAROV

Rechtstranslatorische Verfahren der Künstlichen Intelligenz. Frame-semantische Überlegungen zur Übertragung kulturspezifischer Rechtsbegriffe bei der KI-generierten Urkundenübersetzung	93
---	-----------

FILIP KALAŠ / PAVOL LIPTÁK

Bewertung von Sprachmodellen bei der Übersetzung metaphorischen Denkens in der Wirtschaftssprache	109
--	------------

PAVOL LIPTÁK / FILIP KALAŠ

Sprache als Verkaufsinstrument: Fahnenwörter im Marketing und ihre Auswirkungen auf den Verkauf im Zeitalter der KI	125
--	------------

ISABELLA FERRON

KI-Einsatz zur Übersetzung deutsch-italienischer Wirtschaftstexte: Chance oder Risiko?	141
---	------------

RUI MA

When AI Tastes Iberian Ham: Cultural Gaps in Spanish- Chinese Machine Translation of Gourmet Tourism Products	167
--	------------

SARAH THEROINE / LAURENT GAUTIER /

CHRISTOPHE CRUZ / GENNARO ASCIONE

Textual Molds for In-Context Learning at Sea: Can Generative AI Produce Shipping Forecasts without Semantic Guidance?	185
--	------------

INES-A. BUSCH-LAUER

KI-Berichterstattung zwischen Euphorie und Realität – KI-Newsletter als Informationsmedium	203
---	------------

FACHSPRACHENVERMITTLUNG MIT KI-UNTERSTÜTZUNG

YVONNE HÖRMANN

KI als Schlüssel zur Fachsprache – Personalisierte Lernwege und Kollaboration im Fachunterricht	225
--	------------

SANDRA BRAUN

Präsentationen im Zeitalter Künstlicher Intelligenz: Einsatz, Reflexion und Beurteilung im fremd- und fachsprachlichen Unterricht an Hochschulen	239
---	------------

JANINA M. VERNAL SCHMIDT

Auf dem Weg zu einem kritisch-engagierten Ansatz und einem KI-integrierenden Curriculum im BA-Studiengang „Wirtschaftskommunikation Deutsch“	257
---	------------

ANJA LANGE / GULDASTAN ISMAILOVA

Fachdeutsch vermitteln mit KI: Didaktisierung von Lehrmaterialien und neue didaktische Konzepte durch den Einsatz von Künstlicher Intelligenz	273
--	------------

FACHÜBERSETZEN UND DOLMETSCHEN MIT KI-UNTERSTÜTZUNG

JOSÉ TOMÁS CONDE RUANO

Satisfaction with the Use of Microsoft 365 Copilot in Technical Translation Training	293
---	------------

CARLOS J. GUERRERO RAMOS / ANJA ANNA KÄUPER Entre el entusiasmo y el escepticismo: El uso de la IA generativa para la mejora de la práctica docente y la calidad en la traducción de la divulgación científica	307
ANJA LANGE / IRYNA GAMAN Moderner Übersetzungsunterricht unter Nutzung digitaler Technologien	339
ELISA SERNA-MARTÍNEZ Loosening the Tightrope: Mindfulness, Phenomenology of Language and Embodied Listening in Interpreter Training	361
JESSICA PÉREZ-LUZARDO DÍAZ / BENEHARO ÁLVAREZ-PÉREZ New Trends in the Interpreter's Toolbox: AI's Impact on Pre-Interpreting Tasks	379
E. MACARENA PRADAS MACÍAS / EMILIA IGLESIAS FERNÁNDEZ Ist KI überzeugend? Dolmetschen in Weinverhandlungssituationen	399
JIE LI / PEIYU LIANG / DEJUN LIU Paradigmenwechsel für Medientexte im Zeitalter der KI: Kopplungsanalyse von Nutzern, Themen und Affekten	419
Epilog: Thesen zur Perspektive von KI in der Fachkommunikationsforschung und -lehre	441
Über die Autorinnen und Autoren / About the Authors	445

Vorwort

Wenn KI auf Fach und Sprache trifft – Zur Einführung in den Band

Am 11. Dezember 2015 – also vor ziemlich genau 10 Jahren – beschleunigte sich mit der Gründung des Unternehmens OpenAI die rasante Entwicklung von Künstlicher Intelligenz (KI), die mit der Veröffentlichung von ChatGPT und weiterer generativer Sprachmodelle zu einer ungeahnten Dynamik der KI-Ära geführt hat. Schon heute verändert KI unser Leben stark. Und wir stehen erst am Anfang einer Entwicklung, die in Zukunft in nahezu alle Bereiche des menschlichen Lebens nachhaltig eingreifen wird.

Einerseits bieten KI-Technologien vielfältige Einsatzmöglichkeiten, z. B. zur schnellen Recherche von Informationen, zur Erstellung und/oder Korrektur von Dokumenten, zur Bearbeitung und Erklärung von Aufgaben im Bildungssektor oder zur Optimierung und Personalisierung von Informationsströmen. Andererseits wächst die Skepsis gegenüber KI-generierten Inhalten, vor allem vor dem Hintergrund von Informationsmanipulation und -verzerrung, „Halluzinationen“ sowie unklarer KI-Nutzungsbedingungen. Diese Polarisierung zur Wahrnehmung von KI zeigt sich auch deutlich im Bildungsbereich. Während Lernende die neuen Technologien oft (ohne Faktencheck) nutzen, verzichten Dozierende oft noch immer (aufgrund von Vorbehalten) darauf, sich der Tragweite von KI in Forschung und Lehre bewusst zu werden. Schmidt (2024) hat diese Ambivalenz sehr klar umrissen: „Der Schlüssel liegt darin, KI sinnvoll zu nutzen, ihre Grenzen zu kennen und sie als das zu sehen, was sie ist: eine Technologie, die uns unterstützen soll – nicht ersetzen.“

Die rasante Entwicklung von KI verändert auch die Angewandte Linguistik in nahezu allen Bereichen wie Susoy/Zarfsaz (2025) in ihrer bibliometrischen Studie zu KI-Publikationen dargestellt haben. Ob in den Medien, der Fachkommunikationsforschung, der Fachsprachenvermittlung oder der Fachübersetzung – KI bietet neue Möglichkeiten, stellt uns aber auch vor komplexe

Herausforderungen. Zeit für eine kurze Bilanz zum Status quo im Jahr 2025. In diesem Sammelband möchten wir die Auswirkungen von KI auf die Fachkommunikation aus verschiedenen Perspektiven beleuchten, Projekte vorstellen und zukunftsweisende Thesen diskutieren. Der Fokus der 22 Beiträge von insgesamt 32 Autorinnen und Autoren aus neun Ländern liegt auf vier Bereichen:

1. *Fachkommunikation und KI – Modelle, Konzepte & Terminologie* (Beiträge von ROELCKE, MUNDORF, WISSIK sowie SUCHOWOLEC/LANG),
2. *Fachkommunikationsforschung mit und zu KI* (Beiträge von NAZAROV, KALAŠ/LIPTÁK, LIPTÁK/KALAŠ, FERRON, MA, THEROINE/GAUTIER/CRUZ/ASCIONE, BUSCH-LAUER sowie LI/LIANG/LIU),
3. *Fachsprachenvermittlung mit KI-Unterstützung* (Beiträge von HÖRMANN, BRAUN, VERNAL SCHMIDT sowie LANGE/ISMAILOVA),
4. *Fachübersetzen und Dolmetschen mit KI-Unterstützung* (Beiträge CONDO RUANO, GUERRERO RAMOS/KÄUPER, LANGE/GAMAN, SERNA-MARTINEZ, PÉREZ-LUZARDO DÍAZ/ÁLVAREZ-PÉREZ sowie PRADAS MACÍAS/IGLESIAS FERNÁNDEZ).

In den vier dargestellten Kategorien wird die Vielfalt der Optionen zur Anwendung von KI in der Fachkommunikation exemplarisch veranschaulicht. Es wird deutlich, dass die intensive Erprobung von KI-Modellen und das Experimentieren mit KI-Systemen zu einer neuen Stufe der Auseinandersetzung mit Sprachdynamik im Bereich der Fachsprachen und Fachkommunikation beitragen kann, wobei der Prozess erst in den Anfängen steckt und es nach dieser Phase einer Systematisierung bedarf.

Ich danke allen Autorinnen und Autoren an dieser Stelle herzlich für ihre Bereitschaft, ihre Projekte und Forschungsergebnisse einem breiten Publikum im Open-Access-Format vorzustellen und damit zum Gelingen des Sammelbandes beizutragen. Schließlich möchten wir uns bei den Förderern der Publikation bedanken, insbesondere beim Sächsischen Open-Access-Monografienfonds und der Westsächsischen Hochschule Zwickau, die die Veröf-

fentlichung dieses Sammelbandes durch Finanzierungszusagen erst möglich gemacht haben.

Im Namen der Autorinnen und Autoren möchte ich meinen ganz besonderen Dank an Dr. Anja Steinhauer von textfit für das ausgezeichnete Redigieren und Vereinheitlichen der Beiträge des Sammelbandes aussprechen.

Unser herzlicher Dank gilt auch Frau Astrid Matthes vom Frank & Timme Verlag für ihre freundliche Betreuung dieses Bandes sowie den VerlagsmitarbeiterInnen für die zeitgerechte Umsetzung des Buchprojektes.

Leipzig, im Januar 2026

Die Herausgeberin

Literatur

- SCHMIDT, CH. (2024): „KI im Alltag – Zwischen Hype, Mythen, Realität und Zukunft.“ <https://open.de/blog/ki-im-alltag-zwischen-hype-mythen-realitaet-und-zukunft/> (27.1.2025).
- SUSOY, Z. / ZARFSAZ, E. (2025): „Artificial Intelligence in Applied Linguistics Research: A Bibliometric Study Tracing Between the Years of 2017 and 2024.“ *Educational Academic Research* 58: 141–152.

FACHKOMMUNIKATION UND KI –
MODELLE, KONZEPTE
UND TERMINOLOGIE

Fachliche Kommunikation mit Künstlicher Intelligenz

Abstract

Special communication between representatives from science, technology and institutions is increasingly characterised by the use of Artificial Intelligence. Its use is highly complex not only from an information technology perspective, but also from a linguistic and communicative perspective. For this reason, this essay draws on a general model of special communication expanding it to incorporate Artificial Intelligence and taking into account the producers and recipients of subject-specific texts, these texts as such and their co-text, their general and subject-specific context and the corresponding knowledge, as well as the subject-specific sign systems and their media. For the first time, the model provides an overview of the entire phenomenon of special communication using Artificial Intelligence both from a linguistic and communication science perspective. It is fundamental for identifying individual problems and research areas as well as their multiple future interrelation.

Keywords Fachkommunikation, Kommunikationsmodell, Mensch-Maschine-Kommunikation

1 Vorbemerkungen

Die Interaktion zwischen Menschen und Maschinen im Allgemeinen ist letztlich so alt, wie es Maschinen gibt. Demgegenüber nimmt die Kommunikation zwischen Menschen und Maschinen (nicht: zwischen Menschen mit Maschi-

nen) im Besonderen ihren technischen Beginn spätestens mit der Entwicklung Künstlicher Intelligenz vor rund sechzig Jahren und macht vor allem in den letzten beiden Jahrzehnten enorme Fortschritte. Dies gilt nicht zuletzt auch für Fachkommunikation in Forschung und Lehre, Technik und Handwerk sowie Institutionen und Organisationen unter der Beteiligung von Künstlicher Intelligenz. Ein solch fachkommunikativer Einsatz von Künstlicher Intelligenz ist nicht allein aus informationstechnischer Sicht ausgesprochen komplex, sondern auch aus einer sprach- und kommunikationswissenschaftlichen Perspektive. Dieser Umstand macht es erforderlich, fachliche Kommunikation mit Künstlicher Intelligenz auch und gerade im Rahmen der sprach- und kommunikationswissenschaftlichen Modellbildung zu berücksichtigen. Aus diesem Grund wird in dem vorliegenden Aufsatz auf der Basis eines allgemeinen Modells fachlicher Kommunikation unter Beteiligung von menschlichen Personen schrittweise ein Modell der fachlichen Kommunikation mit Künstlicher Intelligenz entwickelt¹. Dieses Modell berücksichtigt die Produzenten und die Rezipienten fachlicher Texte, die Fachtexte und deren Kontext selbst, einen allgemeinen wie einen fachlichen Kontext und das entsprechende Wissen sowie die fachbezogenen Zeichensysteme und deren Medien. Das Ergebnis ist ein eigenes Modell der fachlichen Kommunikation mit Künstlicher Intelligenz, das einen Überblick über das komplexe Gesamtphänomen erlaubt sowie die Verortung einzelner Problem- und Forschungsbereiche unter Beachtung von deren vielfältigen Bezügen untereinander ermöglicht.

2 Fachkommunikation und Künstliche Intelligenz

Das Konzept fachlicher Kommunikation, das den folgenden Überlegungen zugrunde gelegt wird, folgt dem erweiterten Modell fachlicher Kommunikation, das unter Berücksichtigung von entsprechenden fachkommunikativen Funktionen mit jüngeren Arbeiten von Roelcke (2023, 2025) bereits vorliegt.

.....

1 Herzlichen Dank an See-Young Cho und Koray Parlar für wertvolle technische Hinweise.

Dieses Modell, das auf das einfache Modell fachlicher Kommunikation (vgl. zuletzt Roelcke 2020: 11–13) zurückgeht und in einer längeren Tradition von Zeichen- und Kommunikationsmodellen steht (vgl. Nöth 2000), umfasst vier Bereiche, die in seiner grafischen Darstellung (vgl. Abb. 1) jeweils farbig unterlegt sind.

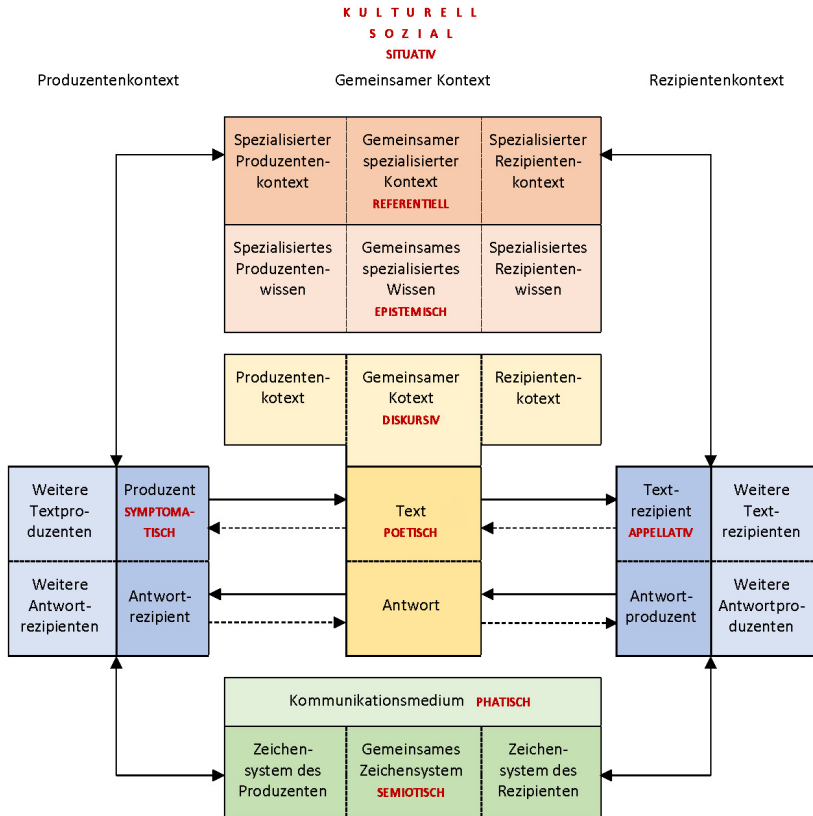


Abb. 1: Modell fachlicher Kommunikation (Roelcke 2023: 119)

In dem ersten Bereich werden der Produzent und der Rezipient eines Fachtextes (als Einzelperson oder als Personengruppe) unterschieden (dunkel- bzw. hellblau unterlegt). Textproduktion und -rezeption selbst werden jeweils als

ein aktiver Prozess aufgefasst (einfache durchgezogene bzw. unterbrochene Pfeile). Dabei ist eine dialogische Struktur vorgesehen, in deren Rahmen Rezipienten von Texten auch als Produzenten einer möglichen Antwort in Erscheinung treten können (diese Struktur ist zudem als vereinfachte Erfassung polylogischer Kommunikation aufzufassen). In funktionaler Hinsicht ist die fachliche Kommunikation einerseits als symptomatisch für den Produzenten und andererseits als appellativ für den Rezipienten anzusehen.

Der fachliche Text und die Antwort auf diesen bilden den zweiten Bereich des vorliegenden Modells (gelb unterlegt); im Hinblick auf die Art und Weise ihrer Gestaltung kommt diesen jeweils eine poetische oder ästhetische Funktion zu. Sie stehen in einem Zusammenhang mit weiteren Texten, welche die Produzenten und Rezipienten jeweils bereits in diesem oder einem anderen Zusammenhang produziert oder rezipiert haben und in deren Kenntnis sie miteinander kommunizieren. Diese Texte bilden jeweils den Kotext des Produzenten und des Rezipienten und weisen letztlich eine diskursive Funktion auf; Texte, die beiden bekannt sind, sind als deren gemeinsamer Kotext anzusehen (hellgelb).

Die fachlichen Texte werden von beteiligten Personen unter Rückgriff auf ein oder mehrere fachlich relevante oder einschlägige Zeichensysteme produziert und rezipiert (grün unterlegt). Hierbei sind verschiedenartige Systeme möglich, deren (gemeinsame) Kenntnis bei dem Produzenten und dem Rezipienten vorauszusetzen ist. Im Rahmen der fachlichen Kommunikation werden solche Systeme durch den Produzenten und den Rezipienten variiert und tradiert, sodass hier ein wechselseitiges Verhältnis (Doppelpfeile) und eine semiotische Funktion anzusetzen sind. Im Weiteren wird ein Zeichenspezifisches Medium der Kommunikation mit phatischer Funktion angenommen (hellgrün).

Unter dem Kontext wird im Modell der außersprachliche Zusammenhang der fachlichen Kommunikation verstanden, auch wenn dieser selbst wiederum in Abhängigkeit von Sprache und Kommunikation steht. In dem vorliegenden Modell werden dabei sowohl ein allgemeiner (gemeinsamer) Kontext (ohne farbliche Unterlegung) als auch ein spezialisierter (gemeinsamer) Kontext (rot unterlegt) angesetzt, die jeweils in einer Wechselbeziehung zu den entsprechenden Produzenten und Rezipienten stehen (Doppelpfeile) und die

referentielle Funktion der Kommunikation ausmachen. Das konzeptionelle und propositionale Wissen der kommunizierenden Personen um und über diesen spezialisierten Kontext ist jeweils durch den Gebrauch von Zeichen bedingt und wird hier mit Blick auf das fachliche Handeln selbst als eigener Teil des Modells unter Beachtung einer epistemischen Funktion angesetzt (hellrot); gemeinsames (spezialisiertes) Wissen ist hier im Sinne von parallelen Kompetenzen und Kenntnissen zu verstehen, die zwischen Produzenten und Rezipienten aufgrund ihrer Kommunikation generiert werden.

Künstliche Intelligenz (KI) oder Artificial Intelligence (AI) wird hier im Allgemeinen als Fähigkeit von Rechenmaschinen bzw. Computern verstanden, Aufgaben zu bewältigen, wie sie im Allgemeinen durch menschliches Denken erledigt werden. Neben einfachen Algorithmen ist hierbei insbesondere an künstliche neuronale Netze, also an komplexe Netzwerke zu denken, die wiederkehrende Muster in großen Datenmengen erkennen und Reaktionen generieren, die solchen von Menschen (mehr oder weniger) entsprechen. Mit Blick auf fachliche Kommunikation wird hier von Sprachmodellen mit großen Textkorpora, sog. *Large Language Models* (LLM), als besonderer Form von Künstlicher Intelligenz ausgegangen, die darauf trainiert werden, menschliche Sprache zu rezipieren und zu produzieren und somit zu interagieren. Solche Sprachmodelle leiten aus einer umfangreichen Menge von Textdaten sprachliche Muster einschließlich entsprechender Textzusammenhänge ab (*self-supervised pre-training*) und sind durch maschinelles Lernen (*Machine Learning*, insbesondere: *Deep Learning*), ein anschließendes überwachtes Feintuning (*supervised fine-tuning*) und bestärkendes Lernen mit menschlichem Feedback (*Reinforcement Learning from Human Feedback*) dazu in der Lage, auf Basis von Wahrscheinlichkeitsverteilungen Textteile oder ganze Texte zu antizipieren und somit auf entsprechende Anfragen zu reagieren (vgl. Blüthgen 2025); so steht etwa *GPT* für *Generative Pre-Trained Transformer*, eine Technologie, in der auf Grund zahlreicher sprachlicher Daten und deren bestehenden Beziehungen neue sprachliche Texte erstellt werden. Auf diese Weise wird eine menschliche Kommunikation simuliert bzw. eine humanoide Kommunikation generiert, die zwar auf Mustererkennung und Generierung basiert, aber nicht auf menschlichem Verstehen oder Bewusstsein.

Neben der Nutzung im Alltag bestehen für eine solche Technologie bzw. entsprechende Systeme zahlreiche professionelle Nutzungsmöglichkeiten (zu einer fachlichen Kommunikation in und mit verschiedenen, darunter auch elektronischen Medien vgl. Roelcke/Breeze/Engberg 2025: 561–743). Hierzu zählen zum Beispiel die professionelle Erstellung von Texten (E-Mails, Beiträge auf Foren, Dokumentationen usw.), die textbasierte Analyse von Daten (Zusammenfassung von Nachrichten oder Unternehmens- bzw. Forschungsberichten), die Unterstützung im Kundenservice (Chatbots, Beantwortung FAQs), der Einsatz in der Sprachmittlung (Übersetzen und Dolmetschen), das Management von Terminen (Logistik), die Durchführung von Recherchen (in Presse oder Wissenschaft) und viele andere mehr. Entsprechende technologische Angebote (die neben Ertel 2025; Gethmann et al. 2022; Russell/Norvig 2021 und einigen Internetquellen wie etwa Wikipedia ebenfalls – wenn auch nur äußerst marginal – bei der Abfassung des vorliegenden Kapitels behilflich waren) bestehen etwa in GPT-4.5 (Open AI), Gemini 2.5 Pro (Google), Claude Opus 4 bzw. Claude Sonnet 4 (Anthropic) und LLama 4 (Meta AI / Facebook).

Die Rezeption von sprachlichen Eingaben durch solche Technologien erfolgt dadurch, dass diese Eingaben zunächst in kleinere Elemente (sog. *Tokens*) aufgegliedert und dann vor dem Hintergrund ihres Zusammenhangs in ein eigenes Beziehungsgefüge gestellt werden. Auf diese Weise wird die Gesamtheit des eingegebenen Textes formal wie funktional (semantisch) selbst re-konstruiert. Die Produktion eigener Texte wird durch die lineare Selektion von Elementen vorgenommen, die den bereits bestehenden Kombinationen in dem zugrundeliegenden Material an Texten in Auswahl und Häufigkeit am meisten entsprechen und somit eine (für den Menschen) relevante und kohäsive (wie kohärente) Replik bilden. Die Herausforderung für die Modellierung der fachlichen Kommunikation zwischen Mensch und Maschine besteht daher insbesondere darin, dass durch solche Systeme menschliche Interaktion nur unter Rückgriff auf umfangreiche Text- bzw. Datenmengen nachgeahmt wird und somit wesentliche Elemente, Relationen und Funktionen der menschlichen Kommunikation spezialisierter Tätigkeitsbereiche außer Acht gelassen werden.

3 Fachkommunikation mit Künstlicher Intelligenz

Die Modellierung von Fachkommunikation erfolgt hier in vier Schritten, die den vier Bereichen des allgemeinen Modells fachlicher Kommunikation entsprechen: Produzenten und Rezipienten (blau), Text und Kotext (gelb), Kontext und Wissen (rot) sowie Zeichensystem und Medium (grün). Sie erfolgt dabei durch entsprechende Änderungen in dem Modell, die im Folgenden Schritt für Schritt eingeführt und erläutert werden (vgl. Abb. 2).

Hinsichtlich des oder der menschlichen Produzenten und menschlichen Rezipienten, die auf der linken Seite des Modells als potentielle Initiatoren der fachlichen Kommunikation mit Künstlicher Intelligenz lokalisiert werden, sind im Großen und Ganzen keine oder kaum Änderungen vorzunehmen: Die betreffende(n) Person(en) kommuniziert/kommunizieren symptomatisch (nicht zuletzt auch mit Blick auf deren Einstellung zu Künstlicher Intelligenz) und antizipieren dabei mehr oder weniger die Technologie, mit der sie es im Rahmen der konkreten Kommunikation zu tun haben (zum Beispiel durch sog. *Prompting*, eine Technik, gezielt Anweisungen oder Fragen zu formulieren, um eine bestimmte Reaktion der Künstlichen Intelligenz hervorzurufen). Die Funktion, die mit den entsprechenden Textäußerungen verbunden ist, besteht in einer Anweisung, die gegenüber der Künstlichen Intelligenz erfolgt, und wird an dieser Stelle somit nicht als appellativ, sondern als instruktiv charakterisiert. Das System der Künstlichen Intelligenz, das an der fachlichen Kommunikation beteiligt ist, handelt nicht im Sinne einer (kreativen, innovativen, eigenständigen, bewussten usw.) menschlichen Tätigkeit, sondern verarbeitet lediglich Information, und wird in der Abbildung daher nicht blau, sondern grau unterlegt. Symptomatisch ist dessen Informationsverarbeitung mit Blick auf die Spezifika der Technologie, die jeweils zum Einsatz kommt. Eine appellative Funktion ist mit dieser Informationsverarbeitung nicht direkt, sondern insofern indirekt verbunden, als die generierten Texte von dem oder den menschlichen Rezipienten als kommunikativ relevant akzeptiert werden. Eine Ausnahme hiervon besteht in dem Fall, dass der Einsatz von Künstlicher Intelligenz bewusst durch eine oder mehrere menschliche Personen erfolgt (in

Entsprechung hierzu kann auch eine mittelbare Antizipation des menschlichen Gegenübers künstlicher Intelligenz angenommen werden).

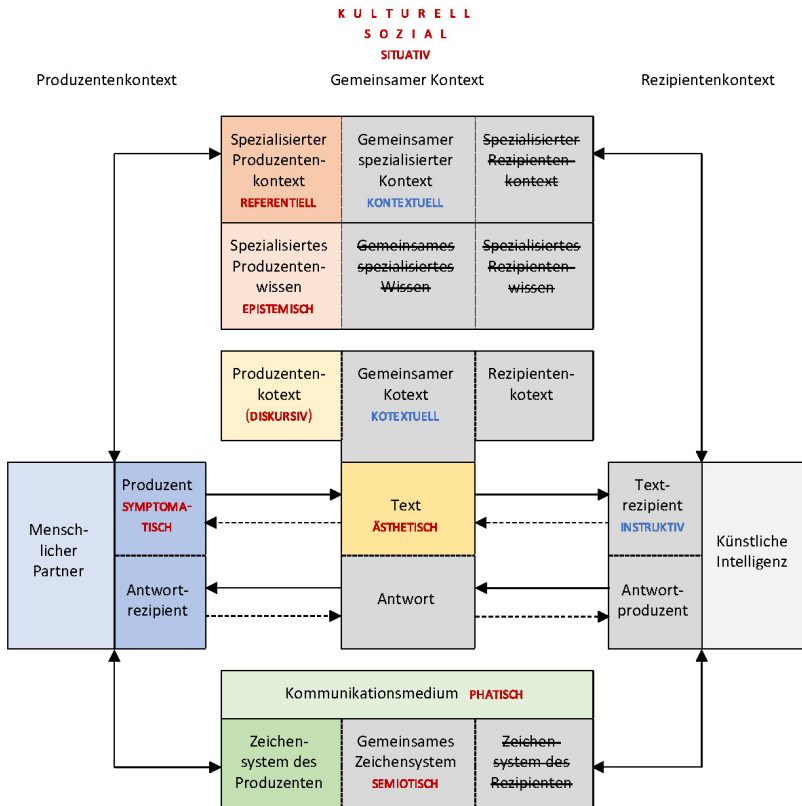


Abb. 2: Modell fachlicher Kommunikation mit Künstlicher Intelligenz

Der Text, der von einem menschlichen Produzenten an das KI-System gerichtet wird, entspricht dessen persönlicher Kompetenz sowie der (mehr oder weniger) antizipierten Verarbeitung von sprachlicher Information durch das entsprechende System. Der Text, der von einer KI generiert an das menschliche Gegenüber gerichtet wird, ist das Ergebnis der technischen Verarbeitung von Daten, die auf der Grundlage von großen Textkorpora sowie einschlägigen

Textrezeptions- und -produktionsprogrammen erfolgt und dabei nicht zuletzt auch von der Art und Weise des *Promptings*. Dieser Text entspricht somit jeweils der (technischen) Kapazität des Datenverarbeitungssystems und wird im Modell daher nicht gelb, sondern grau unterlegt. In beiden Fällen ist eine textuelle Selbstreferenz gegeben, sodass diesen Texten jeweils eine eigene poetische oder ästhetische Funktion zukommt. In diesem Sinne beruht Ästhetik auf menschlicher Kreativität oder auf maschineller Verarbeitung und besteht in der Gestaltung der betreffenden Texte selbst.

Im Bereich des Kotexts sind zwei Interpretationen möglich: Wird der Kotext schlicht als die Menge der (auf welche Art und Weise auch immer) bekannten und der Kommunikation zugrundeliegenden Texte angesehen (was mit Blick auf das erweiterte Modell fachlicher Kommunikation zu kurz greift), ist keine Änderung im vorliegenden Modell erforderlich. Wenn der Kotext jedoch darüber hinaus auch als Kenntnis von Texten einschließlich ihrer Produktions- wie Rezeptionsbedingungen und -verhältnisse betrachtet wird (so, wie dies im erweiterten Modell fachlicher Kommunikation tatsächlich auch geschieht), ist es notwendig, diesen Kotext allein im Bereich der KI auf die Gesamtmenge des sprachlichen Materials zu reduzieren und in dem Modell entsprechend auszugrauen. Entsprechendes gilt dann auch für den gemeinsamen Kotext, der hier letztlich ebenfalls auf das gemeinsame sprachliche Material reduziert wird. Im beiden Fällen verliert der Kotext in weiten Teilen seine diskursive Funktion, da er diese ausschließlich gegenüber dem menschlichen und nicht gegenüber dem maschinellen Partner der Kommunikation entfaltet. Somit ist an dieser Stelle des Modells alternativ eine kotextuelle Funktion im Sinne einer (eigenen bzw. gemeinsamen) konstitutiven Grundlage des betreffenden Textes anzusetzen, nicht zuletzt auch deshalb, weil die Textbeiträge des menschlichen Partners im Falle bestimmter Technologien in die zugrundeliegende Textmenge der KI Eingang finden, wobei die diskursive Textfunktion ausschließlich dem menschlichen Kommunikationspartner vorbehalten bleibt.

Mit Blick auf Kontext und Wissen sind für den menschlichen Kommunikationspartner ein situativer, ein sozialer und ein kultureller Kontext anzusetzen, die nicht zuletzt auch die Art und Weise des Umgangs mit der KI-Technologie betreffen. Die Künstliche Intelligenz kommt zwar ihrerseits in verschiedenen Kontexten zum Einsatz, diese dürfen jedoch mit Blick auf die Funktion der

Datenverarbeitung vernachlässigt werden; dabei ist indessen zwischen Modellen zu unterscheiden, die allein auf bereitgestellte bzw. antrainierte Daten zurückgreifen, und solchen, die durch Mechanismen wie *Retrieval-Augmented Generation* (RAG) auf externe und aktuelle Information zurückgreifen.

Während der menschliche Kommunikationspartner im Rahmen einer spezialisierten Tätigkeit einen bestimmten Ausschnitt der Wirklichkeit fachlich konstituiert und somit einen spezialisierten Kontext schafft, über den er ein entsprechendes Wissen generiert, ist dies im Rahmen der Verarbeitung von (sprachlichen) Daten durch die KI-Technologie nicht der Fall: KI schafft durch ihre Anwendung keinen eigenen Kontext und kein eigenes Wissen, sondern suggeriert dieses nur (mit zunehmender Spezialisierung aufgrund der engeren Datenlänge jeweils umso schwächer). Somit müssen die entsprechenden Elemente aus dem Modell herausgenommen werden. In diesem Kontext ist auf „Reasoning-Modelle“ hinzuweisen, die Ansätze zur Nachvollziehbarkeit des Vorgehens von Künstlicher Intelligenz bieten, auch oder gerade weil diese eben nicht zwischen wahr und falsch zu unterscheiden vermag. Da der spezialisierte Kontext des menschlichen Kommunikationspartners durch den Einsatz von KI beeinflusst wird, kann letztlich so etwas wie ein gemeinsamer spezialisierter Kontext angenommen werden, ein gemeinsames spezialisiertes Wissen im Sinne von wie auch immer vergleichbaren Kenntnissen und Kompetenzen hingegen nicht.

In der grafischen Darstellung des Modells werden daher der spezialisierte Kontext und das spezialisierte Wissen des menschlichen Kommunikationspartners beibehalten. Der gemeinsame spezialisierte Kontext wird angesichts der Beeinflussung des menschlichen Kontexts durch die KI lediglich ausgegraut, während der spezialisierte Kontext, das spezialisierte Wissen und das gemeinsame spezialisierte Wissen in diesem Falle ganz gestrichen werden; hier wie in den analogen Fällen im Folgenden könnten die entsprechenden Felder auch ohne graue Unterlegung erscheinen. Die referentielle und die epistemische Funktion werden ausschließlich dem spezialisierten Kontext bzw. dem spezialisierten Wissen des menschlichen Kommunikationspartners zugewiesen, dem gemeinsamen spezialisierten Kontext zumindest eine kontextuelle Funktion.

Das Interaktionsmedium, über das Mensch und KI-Technologie kommunizieren wird, ist elektronisch. Ganz gleich, ob Texteingaben schriftlich oder mündlich erfolgen, ist eine phatische Funktion im Rahmen der Kommunikation zwischen Mensch und Maschine gegeben. Während Menschen im Rahmen sprachlicher Kommunikation auf kognitiv repräsentierte und linguistisch reformulierte sprachliche Systeme zurückgreifen, erfolgt die Kommunikation über KI durch einschlägige Textrezeptions- und -produktionsprogramme, die mit Quantitäten und Mustern, nicht aber mit einem (lexikalischen) Inventar und (grammatischen) Regeln arbeiten. Aus diesem Grund wird im vorliegenden Modell zwar an dem Zeichensystem des (menschlichen) Produzenten festgehalten, dasjenige des (künstlichen) Rezipienten jedoch aufgegeben und in der grafischen Darstellung gestrichen. Da die KI-generierten Textbeiträge zur Entwicklung und Festigung des sprachlichen Systems der menschlichen Kommunikationspartner beitragen (können) und somit indirekt ebenfalls eine semiotische Funktion innerhalb der fachlichen Kommunikation erfüllen, wird in der grafischen Fassung des Modells ein gemeinsames Zeichensystem in grauer Unterlegung angesetzt.

Ein abschließender Blick auf die grafische Darstellung des Modells fachlicher Kommunikation mit KI zeigt auf der Seite der menschlichen Kommunikationspartner als Initiatoren/Produzenten der betreffenden Kommunikation kaum Änderungen. Die wesentliche Änderung besteht letztlich in der Antizipation und Reflexion der kommunikativen Kapazität der KI-Technologie selbst, wobei diese bei verschiedenen Personen sehr unterschiedlich ausfallen mag. Auf der Seite der KI, die im vorliegenden Fall zunächst mit dem Rezipienten fachlicher Kommunikation verbunden wird, wird von einer maschinellen Rezeption und Produktion von Texten im Rahmen von Textrezeptions- und -produktionsprogrammen ausgegangen, sodass das Feld des Rezipienten und das von dessen Kontext ausgegraut erscheinen. Der Kontext, insbesondere der spezialisierte Kontext und das hiermit verbundene spezialisierte Wissen sowie das Zeichensystem des Rezipienten werden vor diesem Hintergrund in dem grafischen Modell nicht nur ausgegraut, sondern darüber hinaus gestrichen (die entsprechenden Felder könnten somit auch ohne Unterlegung erscheinen). Auch wenn die Kommunikation durch Künstliche Intelligenz anderen Prinzipien folgt als die menschlicher Partner, können neben der automatischen

Antwort sowohl ein gemeinsamer Kontext als auch (hinsichtlich der Referenz für den beteiligten Menschen) ein gemeinsamer Kontext sowie (mit Blick auf die Konstitution eines Zeichensystems für den Menschen) ein gemeinsames Zeichensystem angenommen werden. Von dem Ansatz eines gemeinsamen spezialisierten Wissens im Sinne von vergleichbaren kognitiv fundierten Kenntnissen und Kompetenzen wird demgegenüber abgesehen.

4 Abschließende Bemerkungen

Die Entwicklung des vorliegenden Modells fachlicher Kommunikation mit Künstlicher Intelligenz erfolgt auf Grundlage des bereits bestehenden Modells fachlicher Kommunikation. Angesichts der zahlreichen Bereiche, die in dem Modell der Fachkommunikation mit KI gegenüber dem allgemeinen Modell fachlicher Kommunikation geändert oder ausgespart werden, stellt sich abschließend die Frage, ob diese Vorgehensweise letztlich überhaupt sinnvoll ist und nicht vielmehr eine eigene Modellierung angemessen wäre. Ohne einer solch spezifischen Modellierung vorzugreifen oder deren Zweckhaftigkeit infrage zu stellen, sprechen mindestens zwei Gründe für die hier gewählte Vorgehensweise: Zum einen werden in diesem Modell die Gemeinsamkeiten und Unterschiede beider Arten von fachlicher Kommunikation herausgearbeitet, und zum anderen bleibt die humanzentrierte Perspektive auf die Fachkommunikation in Forschung und Lehre, Technik und Handwerk sowie Institutionen und Organisationen erhalten; eine Perspektivierung, welche Künstliche Intelligenz selbst zum Ausgangspunkt nimmt, sollte somit eher einen technologischen Ansatz wählen. Im Idealfall werden diese beiden Perspektiven zusammengeführt, um eine Optimierung fachlicher Kommunikation mit Künstlicher Intelligenz zu erreichen, die sowohl dem Menschen als auch der Maschine als Partner gerecht wird. Dies gilt sowohl für den Forschungs- und Anwendungsbereich im Ganzen als auch für dessen Teile sowie deren Bezüge im Einzelnen. Das Ergebnis wäre ein kombiniertes Modell der fachlichen Kommunikation von menschlichen Personen mit KI-Technologie, das letztlich einen polyperspektivischen Überblick über das komplexe Gesamt-

phänomen erlaubt sowie eine noch bessere Verortung einzelner Problem- und Forschungsbereiche unter Berücksichtigung von deren vielfältigen Bezügen untereinander.

Literatur

- BLÜTHGEN, CHRISTIAN (2025): „Technische Grundlagen großer Sprachmodelle.“ *Die Radiologie* 4: 227–234. <https://www.springermedizin.de/kuenstliche-intelligenz/radiologie/technische-grundlagen-grosser-sprachmodelle/50731384> (25.07.2025).
- ERTEL, WOLFGANG (2025): *Grundkurs Künstliche Intelligenz. Eine praxisorientierte Einführung*. 6. Auflage. Wiesbaden: Springer Vieweg.
- GETHMANN, CARL FRIEDRICH / BUXMANN, PETER / DISTELRATH, JULIA / HUMM, BERNHARD G. / LINGNER, STEPHAN / NITSCH, VERENA / SCHMIDT, JAN C. / SPIEKER, GEN. DÖHMANN, INDRA, HRSG. (2022): *Künstliche Intelligenz in der Forschung. Neue Möglichkeiten und Herausforderungen für die Wissenschaft*. Wiesbaden: Springer.
- NÖTH, WINFRIED (2000): *Handbuch der Semiotik*. 2., vollständig neu bearbeitete und erweiterte Auflage. Stuttgart, Weimar: Metzler.
- ROELCKE, THORSTEN (2020): *Fachsprachen*. 4., neu bearbeitete und wesentlich erweiterte Auflage. Berlin: Schmidt.
- ROELCKE, THORSTEN (2023): „Funktionen fachlicher Kommunikation.“ *Glottology* 14.1: 111–124.
- ROELCKE, THORSTEN (2025): „Models and Concepts of Specialized Communication.“ *Specialized Communication. An International Handbook*. Volume 1. Eds. Thorsten Roelcke / Ruth Breeze / Jan Engberg. Berlin, Boston: De Gruyter Mouton. 3–29.
- ROELCKE, THORSTEN / BREEZE, RUTH / ENGBERG, JAN, EDS. (2025): *Specialized Communication. An International Handbook*. Volume 1. Berlin, Boston: De Gruyter Mouton (Handbooks of Linguistics and Communication Science / Handbücher zur Sprach- und Kommunikationswissenschaft, HSK 47.1).
- RUSSELL, STUART / NORVIG, PETER (2021): *Artificial Intelligence. A Modern Approach*. 4th ed. London: Pearson.

Rechtskommunikation im Umbruch

Ein Framework für *AI Legal Literacy* und fachkommunikative Kompetenzen im Recht

Abstract

Artificial Intelligence (AI) and algorithmic systems are transforming law from a text-based domain into one shaped by data and algorithms. Although legal experts can expand their level of professional specialization through AI, they risk losing their expertise if they lack the ability to use, manage, and govern AI responsibly. This profound transformation raises questions about evolving practices in specialized communication: among experts within a field and from different disciplines, between experts and laypersons, and in human-machine interaction across technical, institutional, legal, and ethical contexts. It also affects socio-technical and socio-linguistic aspects, as well as specialized communication and competence requirements. The concept of *AI literacy* has so far been a desideratum in specialized communication research. This article therefore develops and operationalizes a competence compass for AI legal literacy (AILL) focusing on specialized communication at (1) horizontal, (2) vertical, (3) systemic, and (4) functional levels, and (5) the discursive intersection, considering attitudes, behavior and knowledge.

Keywords AI Literacy, KI-Kompetenz, language awareness, legal linguistics, Rechtskommunikation

1 Einleitung: Quo vadis, Rechtskommunikation, im Zeitalter von KI?

In einem familiengerichtlichen Beschluss des Amtsgerichts Köln heißt es in der schriftlichen Begründung: „Die weiteren von dem Antragsgegnervertreter im Schriftsatz vom 30.06.2025 genannten Voraussetzungen stammen nicht aus der zitieren [sic!] Entscheidung und sind offenbar mittels künstlicher Intelligenz generiert und frei erfunden.“¹ [Herv. i. Orig. – MM] Es folgt eine Aufzählung von „Halluzinationen“ aus dem Schriftsatz des Rechtsanwalts, wie sie typischerweise beim Einsatz von generativer KI auftreten können: erfundene Rechtsquellen, nicht-existente oder vermischte Fundstellen, falsche Autoren. Der Richter unterstellt dem Anwalt bewusstes Verbreiten von Unwahrheiten und Verfälschung der Rechtslage und weitreichende Folgen: „Der Verfahrensbevollmächtigte hat derartige Ausführungen für die Zukunft zu unterlassen, da sie die Rechtsfindung erschweren, den unkundigen Leser in die Irre führen und das Ansehen des Rechtsstaates und insbesondere der Anwaltschaft empfindlich schädigen.“ Unabhängig davon, ob man die Fehler in dem anwaltlichen Schriftsatz als Täuschungsversuch oder als schlampige fachliche Arbeit interpretiert, wird man dem Anwalt kaum fachliche *AI Literacy* zugestehen; dem Richter vermutlich durchaus, schon deshalb, weil er die angegebenen Quellen sorgfältig prüfte und die Fehlerursache hier zu Recht auf unsachgemäße KI-Nutzung zurückführte. ‚Betroffen‘ sind neben dem Verfasser (Anwalt) und dem Leser (Richter) direkt und indirekt: Prozessbeteiligte, die Anwaltschaft als Berufsstand, die Richterschaft sowie die gesamte Öffentlichkeit. *AI Literacy* ist also in einer Zeit, in der KI-Technologien längst Teil der Rechtskommunikation geworden sind, keine Marginalie – weder in der Justiz noch in sämtlichen Gesellschaftsbereichen, die von Recht umgeben und durchdrungen sind (vgl. Felder/Vogel 2017: XVIII).

.....
1 S. AG Köln, Az. 312 F 130/25, Beschluss vom 02.07.2025.

Unter dem Einfluss von Künstlicher Intelligenz (KI) und algorithmischen Systemen² (*Algorithmic-Decision-Making-Systems*, Krüger/Lischka 2018: 9) verändert sich die Domäne des Rechts rasant und fundamental. Der Domänenbegriff wird hier verstanden als Wissensdomäne mit den korrespondierenden Handlungsfeldern (Felder 2024, 2008) sowie als soziale Entität mit indifferenten sprachlich-kommunikativen Praktiken und dementsprechend divergierenden „Konventionen, Werten und Erwartungen an die Art und Weise sprachlichen Handelns“ (Jakobs 2023: 60). KI-Technologien haben als „disruptive Technologien“ (Bower/Christensen 1995) das Potenzial, Produkte, Dienstleistungen und Prozessketten zu verändern, zu verbessern, weiterzuentwickeln und komplett zu ersetzen. Besonders stark beeinflussen generative KI und große Sprachmodelle (*Large Language Models*, LLMs) die Wissensarbeit (Butollo et al. 2024; Meier 2025: 15–23; Tomlinson et al. 2025): Im Gegensatz zu früheren technologischen Innovationswellen sind Hochqualifizierte und Wissensarbeiter:innen besonders betroffen, denn das „Substituierbarkeitspotenzial“ steigt mit dem Anforderungsniveau beruflicher Tätigkeiten (Grienberger et al. 2024: 7). Generative KI versetzt Expert:innen eher in die Lage, Aufgaben auch jenseits ihrer Kernexpertise erfolgreich zu bearbeiten (Sack et al. 2024). Demnach können Fachpersonen ihren Fachlichkeits- und Spezialisierungsgrad einerseits ausdifferenzieren und erweitern; andererseits bedrohen die technologischen Möglichkeiten ihre (alleinige) Expertenschaft. Immer mehr Beschäftigte setzen KI-Technologien am Arbeitsplatz täglich ein (Hartley et al. 2025: 18), wobei die formelle Einführung von KI auf Unternehmensebene der tatsächlichen Nutzung noch hinterherhinkt (Arntz et al. 2025: 18). Ähnliches lässt sich auch für den Bereich der Justiz feststellen (Gutiérrez 2024).

Aus diesen Entwicklungen ergeben sich drängende Fragen für die Fachkommunikation: nach veränderten Praktiken in der Interaktion von Mensch und Maschine, nach der Kommunikation zwischen Expert:innen innerhalb eines Fachs und aus verschiedenen Disziplinen (horizontal) sowie zwischen Expert:innen und Lai:innen (vertikal); zu den technischen, institutionellen,

.....
 2 Zum Begriff vgl. Dickert/Otte (2022). S. auch iRights.Lab (2025: 26) in Anlehnung an die KI-Verordnung der Europäischen Union und den „Umsetzungsleitfaden zur KI-Verordnung“ der Bitkom e.V. (2024).

rechtlichen und ethischen Rahmungen; nach soziotechnischen, -linguistischen und fachkommunikationswissenschaftlichen Beschreibungsmöglichkeiten und der Frage nach den domänenspezifischen Kompetenzanforderungen in der Fachkommunikation mit (fach-)didaktischen Implikationen.

AI Literacy nimmt nicht zuletzt als Folge der gesetzlichen Pflichten breiten Raum im ökonomischen, rechtlichen und politischen Diskurs ein, wird von unterschiedlichen Forschungsrichtungen beforscht und ist schon längere Zeit Gegenstand der (Fach-)Didaktik. In der Fach- und Professionskommunikation werden Kompetenzen in einzelnen Aspekten z. B. als berufsbezogene „Registerkompetenz“ (Efing 2024), als „Kompetenzsignalisierung“ (Rhein 2015: 445f.), umfassender als „Leadership-Kompetenzen“ (Zerfaß et al. 2021: 60) und „Komplexitätskompetenz“ (Busch/Luttermann 2024: 8) erörtert. Das Konzept *AI Literacy* ist in der Fachkommunikationsforschung bislang ein Desiderat und wird beispielsweise im Handbuch *Specialized Communication* (Roelcke et al. 2025) nicht explizit thematisiert. In dem lückenhaft untersuchten Feld erweist sich die Domäne des Rechts als vielversprechend, gilt die Rechtssprache doch zum einen als prototypisch fachsprachliche Varietät (Busse 1999: 1382) und eine „der wirkmächtigsten Fachsprachen in unserer Gesellschaft“ (Wimmer 2008: 81); zum anderen haben KI-Sprachmodelle gerade in rechtlichen Handlungsfeldern hohes Automatisierungspotenzial (McKinsey 2023: 41; Hatzius et al. 2023: 7) und führen potenziell zu Jobverlusten wie auch zum Jobzuwachs für Anwält:innen (World Economic Forum 2025: 21).

2 Rechtskommunikation vom Text zum Code als Prozess

Seit geraumer Zeit entwickelt sich das Recht von einer textbasierten Domäne („Recht als Text“, Busse 1992) hin zu einem daten- und algorithmenbasierten Recht („Law as Code“, Billa 2024, „Law-as-data“, Chau/Livermore 2024, Schulz/Schmees 2022: 585). Damit verändern sich die kommunikativen Arbeitsmodi und -prozesse fundamental.

2.1 „Recht ist Text“

Traditionell bedeutet Handeln im Recht „Spracharbeit“ (Felder/Vogel 2017: IX) und „Textarbeit“ (Müller 1994), die Arbeit mit und an der Sprache einschließt (Wimmer 2009: 237). Dies legt eine sprachlich-kommunikative und handlungsbezogene Bestimmung von *Fach* „als ein (mehr oder weniger stark) spezialisierter menschlicher Tätigkeitsbereich“ (Roelcke 2020: 14) nahe. Rechtskommunikation setzt also kompetenten Umgang mit Zeichen, juristischer Fachsprache und juristischen Fachtextsorten voraus, die „Elemente institutionellen Handelns“ darstellen (Busse 2000: 663). Gerade die Sprach- und Zeichenhaftigkeit des Rechts ermöglicht die Sprachverarbeitung durch *Natural Language Processing* (NLP). Mit der fortschreitenden Digitalisierung und „Datafizierung“ des Rechts (Wörner et al. 2024: 1) und der „Verdatung“ der juristischen Textarbeit (Vogel 2024: 9) bewirken sie einen Paradigmenwechsel in der rechtlichen Fachkommunikation.

2.2 „Law as Code“

Galt lange Zeit das Faxgerät sinnbildlich für überkommene Kommunikation in der Rechtsdomäne, so haben sich dort inzwischen digitale Technologien, regelbasierte Algorithmen sowie KI-Systeme des Maschinellen Lernens etabliert (Hartung/Gerlach 2025; Meisenbacher et al. 2024) – in der *fachinternen* Rechtskommunikation unter Fachleuten sowie in der *fachexternen* Kommunikation zwischen Expert:innen und Lai:innen. Anwendungsbereiche im Recht (*Legal NLP*) lassen sich nach Aufgaben wie Zusammenfassen, Vorbearbeiten, Klassifizieren, Recherchieren, Texterstellung etc. clustern (Katz et al. 2023, vgl. Hartung/Gerlach 2025; Dickert/Otte 2022). Unter den Sammelbegriff *Legal Tech* werden „Software und Online-Dienste, die juristische Arbeitsprozesse unterstützen oder gänzlich automatisiert durchführen“ gefasst (Martinetz/Maringele 2020: 3). Die Bezeichnung gilt nicht nur den Technologien, sondern impliziert „eine Transformation traditioneller Arbeitsweisen und eine Neugestaltung der Beziehung zwischen Anwalt und Mandant“ (Dietrich 2024: 2093, Fußnote 4). Betroffen sind auch „Arbeitsprozessketten“ im Sinne Schuberts

(2007: 133) mit einer typischen Abfolge von fachkommunikativen Arbeiten, die zumindest in Teilen maschinell gelöst werden können.

2.3 „Mehr als Text“

Recht ist insofern nicht (nur) Text oder Code, sondern vor allem auch „Verfahren und Handeln“ (Kocher 2025: 229). In der Struktur und sprachlichen Form von Rechtstexten, deren Sprach- und Textmuster formelhaft erlernt werden, zeigen sich rechtliche Eigenlogiken und die Suche nach dem Allgemeinen mit dem Anspruch auf normative Geltung. Sprachmodelle sind für solche „vorgestanzten Texte“ (Kocher 2025: 234) per se geeignet, sie können Zeichenfolgen und Texte mit einer „Illusion von Kohärenz“ (Lotze 2016) erzeugen und insofern „intelligible Texturen“ (Schneider 2024) darstellen. In der Technischen Redaktion sind Automatisierungsprozesse schon lange üblich und können als „eine konsequente Verlängerung und Erweiterung bisheriger Praktiken mit neuen Mitteln“ betrachtet werden (Lehnen/Steinhoff 2024: 542). Dasselbe lässt sich zu anwaltlichen Fachkommunikationspraktiken feststellen (vgl. Lilienthal et al. 2023). In der Justiz erfolgt der KI-Einsatz zurückhaltender, doch gibt es auch dort eine stetig wachsende Zahl von KI-Projekten (Mielke 2025).

3 *AI Literacy* und *KI-Kompetenz* im Recht

Es liegt nahe, dass die aufgezeigten Umbrüche weitreichende Implikationen für die Kompetenzanforderungen in rechtlichen Arbeitsfeldern haben (Tu 2024).

3.1 **AI Literacy/KI-Kompetenz in der KI-Verordnung (KI-VO)**

Die Legaldefinition aus Artikel 3 Nr. 56 KI-Verordnung 2024/1689 (KI-VO) bestimmt *KI-Kompetenz* (in der deutschsprachigen Fassung) als *Fähigkeiten, Kenntnisse und Verständnis*, hinsichtlich der Art und Weise des KI-Einsatzes (*sachkundig, verantwortungsvoll, sicher*) sowie der *Bewusstheit* hinsichtlich *Chancen und Risiken*. Nach Art. 4 KI-VO müssen Anbieter und Betreiber von KI-Systemen sicherstellen, dass Mitarbeitende über ein „ausreichendes Maß an KI-Kompetenz“ verfügen. Dass die KI-VO kein einheitliches Kompetenzniveau vorsieht, eröffnet Gestaltungsräume für eine bedarfsgerechte Entwicklung von *AI Literacy*, die von Kontext, Nutzergruppen, Komplexität des KI-Systems und Einsatzzweck abhängt (Block et al. 2025: 11; Bundesnetzagentur 2025). Die Ausrichtung am (fachlichen) Hintergrund, der Erfahrung und dem Anwendungskontext erfordert daher fach- und domänenspezifische *AI-Literacy*-Konzepte.

3.2 **AI Literacy und KI-Kompetenz: Entwicklungslinien und Modellierungen**

Der *Literacy*-Begriff ist eng mit Schriftlichkeit, Schreib- und Lesefähigkeiten verknüpft, wurde aber zunehmend für fach- und wissensspezifische Bereiche erweitert (etwa als *Media, Political, Financial, Climate* oder *Legal Literacy*) und als *Data* und *Digital Literacy* auf multiple und multimodale Formate angewendet (OECD 2019: 25; Gilster 1997: 1; Ferrari 2012; Law et al. 2018). Dies führt dazu, dass Modellierungsansätze von *AI Literacy* (bezogen auf KI-Systeme) sich überlappen mit *Media Literacy* (bezogen auf Kommunikationsformen), *Data Literacy* (Umgang mit Daten, deren Interpretation und Entscheidungen) und *Computation Literacy* (algorithmische und Coding-Anwendungen) (OECD 2025; Hackl et al. 2025: 14).

Nach Long/Magerko (2020: 598) bedeutet *AI Literacy* die kritische Einschätzung, effektive Kommunikation und Kollaboration mit KI in verschiedenen Lebensbereichen und am Arbeitsplatz mithilfe von 16 Kernkompetenzen.

Ng et al. (2021: 4) schlagen die Kompetenzdimensionen (1) Wissen und Verstehen, (2) Anwenden, (3) Evaluieren und (4) ethische Aspekte vor. Frühere Konzeptionen adressieren vor allem Personen mit geringem technologischen Fachlichkeitsgrad (Laupichler et al. 2023: 1) mit Basiskompetenzen für Anwender:innen (Laupichler et al. 2022).

AI-Literacy-Konzeptionen im Bildungsbereich liegen meist die Lernziel-taxonomien (nach Bloom 1972; Anderson/Krathwohl 2001) mit den Kompetenzdimensionen von Wissen (Kognition), Fähigkeiten (Verhalten) und Einstellungen (Werten) zugrunde, sie berücksichtigen zum Teil unterschiedliche Interaktionsebenen, Handlungsrollen und fachliche Aspekte (OECD 2025; Anthropic 2025). Dabei stellt sich die Frage einer geteilten Verantwortung (vgl. Wilder et al. 2021) in einer gemeinsamen Verantwortungskultur (Brommer et al. 2023).

3.3 Fach- und domänenbezogene *AI Literacy*

Für eine fachkommunikative Modellierung von *AI Literacy* ist das Zusammenwirken von individuellen Faktoren, der Motivation und inneren Bereitschaft, sozialen Fähigkeiten sowie Handlungsaspekten (Weinert 2002: 27f.), mit organisationalen und gesellschaftlichen Faktoren relevant. Dies gilt etwa bei der Bewertung und Kommunikation von Risiken oder in der Zusammenarbeit in multidisziplinären KI-Projekten, gerade bei Einsatzzwecken im Bereich der Hochrisiko-KI, die gemäß KI-VO nur mit Konformitätsbewertung erlaubt sind (etwa in der Strafverfolgung, der Rechtspflege und demokratischen Prozessen). Die zunehmende Spezialisierung von KI-Technologien in der Rechtskommunikation erfordert spezifische Formen von *AI Literacy* (Gültekin-Várkonyi 2025). Damit werden fach- und domänenspezifische Beschreibungen von *AI Literacy* im Fach notwendig; Letzteres verstanden als „ein spezialisierter Tätigkeitsbereich von Personen mit entsprechender Expertise“, in dem Kommunikation ein konstitutives Element ist (Roelcke 2023: 117). Eine konzeptionelle Grundlage für fachbezogene *AI Literacy* bietet die holistische Kompetenzbewertungsmatrix von Knoth et al. (2024), die nach (1) generischer, (2) domänenspezifischer und (3) KI-Ethikkompetenz unter-

scheidet und diese jeweils für die Dimensionen von Kognition, Verhalten und Einstellung (als Akronym ABC für **A**ttitudes, **B**ehaviour und **C**ognition) beschreibt.

4 Ein Framework für *AI Legal Literacy*

Zur Modellierung von *AI Legal Literacy* in der juristischen Fachkommunikation werden domänenspezifische Kompetenzmodelle (Knoth et al. 2024; Gültekin-Várkonyi 2025) mithilfe deskriptiver fachkommunikativer Modelle erweitert und spezifiziert (Suchowolec 2024: 42f.). Modelle der *horizontalen Fachsprachengliederung* gehen von einer Unterteilung in Fächer und Disziplinen (mit Überlappungsbereichen) aus (Hoffmann 1985; Fluck 1996: 16f.; Roelcke 2020: 43f.). Diese sind Ergebnis der sprachlichen Kommunikation über fachliche Inhalte (Kalverkämper 1990: 100) und beinhalten Fachwissen von Expert:innen als „kulturspezifisches Wissen“, basierend auf „Erfahrungen mit Gegenständen [...], die bestimmten Personengruppen vorbehalten sind“ (Konerding 2015: 72). Nach *vertikalen Gliederungen* kann man verschiedene Konstellationen in der fachlichen Kommunikation zwischen Rechtsexpert:innen (etwa zwischen Richter:in und Anwalt) bzw. in der „fachbezogenen Kommunikation“ (Göpferich 2008: 11; Heidrich/Schubert 2020: 13), etwa zwischen Anwältin und Mandanten oder zwischen juristischen Lai:innen über fachliche Inhalte, unterscheiden.

Die Modellierung von *Fachkommunikation als Prozess* folgt bei Roelcke (2020) einer bilateralen Perspektive von menschlichem Produzenten und Rezipienten. Im KI-Zeitalter ist dieses Verständnis um die „Mensch-Maschine-Interaktion“ zu erweitern, die sich zur „kommunikativen KI“ entwickelt (Hepp et al. 2022), immer stärker kollaborativ und „koaktiv“ erfolgt (Steinhoff 2023), „hybrid“ (Seufert/Meier 2024: 57f.), „augmentiert“ (als Unterstützung bei komplexen Aufgaben) oder „automatisiert“ (als Substitution menschlicher Tätigkeit) (Meier 2025: 17). Die Modellierung von Fachkommunikation als „*fachkommunikative Arbeitsprozesskette*“ (Schubert 2003, 2007) akzentuiert die einzelnen Handlungsphasen in medial unterschiedlichen Kommunika-

ten. Solche Prozessketten werden durch multimodale Funktionalitäten von generativer KI (*Text-to-Speech*, *Speech-to-Text*) fluiden und effizienter bearbeitbar; KI-Agenten können komplexe Arbeitsabläufe selbsttätig organisieren und mittels geeigneter Werkzeuge ausführen.

Nachfolgend wird der Kompetenzkompass für *AI Legal Literacy* mit vier Teilkompetenzen dargestellt, über die Fachpersonen im Recht verfügen sollten, um KI für die intra-, inter- und extrafachliche Kommunikation zu verstehen, zu evaluieren sowie erfolgreich und verantwortungsvoll einzusetzen. Dementsprechend wird *AI Legal Literacy* (AILL) als fachkommunikative KI-Kompetenz auf (1) horizontaler (2), vertikaler, (3) systemischer und (4) funktionaler Ebene jeweils mit Einstellungen (A), Verhalten (B) und Wissen (C) operationalisiert. In der Schnittmenge wird diskursive AILL (D, Discourse) gleichsam als Kompassnadel dargestellt.

Aus Platzgründen wird das Framework exemplarisch anhand des Eingangsbeispiels erläutert: Der Anwalt hatte offenbar unsachgemäß ein KI-Werkzeug zur juristischen Recherche genutzt bzw. das Rechercheergebnis ungeprüft in seinem Schriftsatz weiterverarbeitet. Die Fachkommunikation fand zwischen juristischen Fachpersonen auf *horizontaler Ebene im Kompetenzfeld (1)* statt. Innerhalb der justiziellen Kommunikationsverfahren benötigt der Rechtsanwalt auf der Ebene der *Attitudes (A)* Rollenbewusstheit seiner organisationalen und funktionalen Rolle mit inhärenten Verantwortlichkeiten und berufsrechtlichen Pflichten. Er muss auf Ebene der *Cognition (C)* die Funktion von KI-Werkzeugen verstehen, um Risiken (z. B. unzuverlässige Outputs) einzuschätzen. Weiterhin benötigt er Spezialwissen zu juristischen KI-Technologien (etwa der DSGVO- und berufsrechtskonformen Rechts-KI *Noxtua* mit dem Workspace Beck-Noxtua auf Basis der Fachdatenbank beck-online) sowie zu Möglichkeiten und Grenzen spezialisierter KI für fachliche Aufgaben. Auf der Verhaltensebene *Behaviour (B)* braucht er handlungsbezogen-kollaborative Strategien bei der intrafachlichen Fachkommunikation; im konkreten Fall mit anderen Fachpersonen wie dem Richter; bei der Interaktion mit KI auf *systemischer Ebene im Kompetenzfeld (3)* etwa durch fachsprachliche Strategien und Nutzung von Ontologien, auf *funktionaler Ebene im Kompetenzfeld (4)* durch fachliches Prompting oder die Steuerung von fachkommunikativen Arbeitsprozessketten.

A: fachliche Rollen Anforderung reflektieren
 B: inter-/transdisziplinär/im Team kollaborieren
 C: spezialisierte Rechts-KI-Systeme kennen, verstehen, einordnen

fach(sprach)liches Wissen

A: Fach- als Wissenskultur in LLMs (mit Datenbegrenzungen, Biases) reflektieren
 B: verlässliche KI-Ko-Produktion steuern
 C: Fachsprache, -konzepte & -logiken in Daten, Prompts, KI-Kommunikaten verstehen

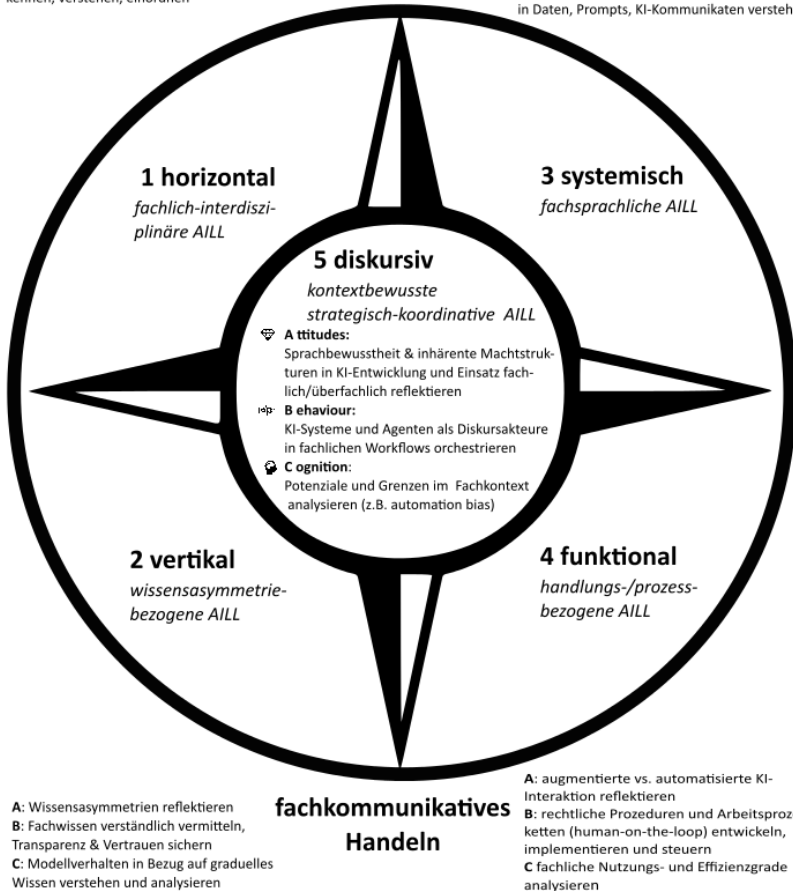


Abb. 1: AI-Legal-Literacy-Kompetenzkompass: Dimensionen und Entwicklungsfelder mit *Attitudes, Behaviour, Cognition, Discourse*

5 Thesen zur Fachkommunikation(sforschung) – (nicht nur) im Recht

Im vorliegenden Beitrag wurde die Transformation der Fachkommunikation in der Domäne des Rechts skizziert und im Anschluss an bestehende Frameworks ein fachkommunikationswissenschaftlich orientierter, domänenspezifischer Rahmen für *AI Literacy* vorgeschlagen, der einerseits auf die (fach-)didaktische Vermittlung und andererseits auf die Forschungsdesiderata der Fachkommunikationswissenschaft abzielt. Die Ergebnisse werden thesenhaft zusammengefasst:

- (1) Durch KI-Technologien verändert sich die juristische Kommunikation *fachintern* (wie das Eingangsbeispiel eindrücklich zeigt) wie *fachextern* (z. B. beim Wissenstransfer in verständlicher Sprache). Auf horizontaler Ebene wird die Notwendigkeit zum verstärkten inner- und interdisziplinären Diskurs erwartet (Hartung/Gerlach 2025: 28), auf vertikaler Ebene eine geringere Nachfrage nach Arbeitskräften für einfache und Routineaufgaben prognostiziert. Spezialwissen, soziale und metakognitive Kompetenzen dürften an Bedeutung gewinnen, einfache kognitive und kommunikative Fertigkeiten hingegen an Relevanz verlieren (Meier 2025: 25). Dies ermöglicht höhere Arbeitsproduktivität und geringere Kosten, gleichzeitig auch einen geringeren Wert (und Wertschätzung) fachlicher Kommunikationsleistungen.
- (2) Die Möglichkeit, komplexe juristische Inhalte verständlich an juristische Lai:innen zu vermitteln, birgt Chancen für verbesserten Zugang zum Recht und eine Verflachung des Wissensgefälles zwischen Fachpersonen und Lai:innen; sie stellt aber auch hohe Anforderungen an KI-Systeme, an Rechtsanwender:innen, die sie verantworten, und Fachpersonen wie Lai:innen, die ihnen vertrauen. *AI Legal Literacy* als fachliche bzw. fachbezogene und auf den Interaktionsprozess bezogene KI-Kompetenz ist unabdingbar, um KI-Systeme verantwortungsvoll für die fachliche Arbeit einzusetzen. Die Transformation der Rechts-

kommunikation betrifft sämtliche Rechts- sowie alle Gesellschaftsbe-
reiche, die von Rechtskommunikation durchdrungen sind. Dies erfor-
dert auch fachbezogene KI-Basiskompetenz für Nicht-Fachpersonen,
die etwa spezialisierte Chatbots für Rechtsfragen nutzen und Fehler
erkennen müssen.

- (3) KI-Technologien verändern das Machtgefüge im Recht und erfordern
eine Verantwortungskultur im Fach und über das Fach hinaus. Sprach-
modelle sind integrativer Bestandteil der Fachkommunikation. Mit ei-
ner Quantifizierung durch Algorithmen gewinnen autonome Systeme
an Wirk- und Entscheidungsmacht im Recht und einer verrechtlichten
Gesellschaft. Fragen der Grenzziehung zwischen einer Augmentation
durch KI und einer Substitution menschlicher Entscheidung sind
weder trivial noch können sie *nur* rechtlich, *nur* technisch oder *nur*
ethisch gelöst werden. Empirische Untersuchungen zu fachkommuni-
kativen Praktiken beim KI-Einsatz könnten bei der Beantwortung
dieser folgenschweren Einschätzungen helfen.
- (4) Fach und Wissenskommunikation im Zeitalter von KI findet zuneh-
mend als Interkommunikation statt und ist als solche zu entwickeln
(vgl. Mundorf/Seidl 2025: 297). Aufgrund der rasant fortschreitenden
technologischen Veränderungen des Rechtsmarkts sind tragfähige
Konzepte für Aus- und Weiterbildung vonnöten. Dazu müssen Formate
etabliert werden, die auf Kontinuität angelegt sind, etwa *Round Tables*
und interdisziplinäre Schulungen, Fachtagungen aus Wissenschaft und
Rechtspraxis, praxisorientierte Hackathons, Agentathons, Foren für
Fachpersonen und Bürger:innen.
- (5) Die Entwicklung und Evaluierung von fachspezifischer *AI Literacy*
benötigt Expertise aus den jeweiligen Fachdomänen und der Fach-
kommunikationsforschung. Teilkompetenzen von *AI Literacy* bauen
sukzessive aufeinander auf. Die Fähigkeit zum kontinuierlichen Lernen
wird dabei als transversale Kompetenz verstanden, die der dynami-
schen technologischen Entwicklung Rechnung trägt (Annapureddy et
al. 2025). *Prompting Skills* werden nach wie vor als Kernkompetenz
eingestuft; mit komplexeren KI-Agenten könnten sich diese Anwender-
fähigkeiten aber hin zum Orchestrieren und Koordinieren von Work-

flows verlagern (McKinsey 2025), was eine stärkere Orientierung an fachkommunikativen Prozessketten mit sich bringt.

Der hier zur Diskussion gestellte Ansatz für die Domäne des Rechts möchte mit dem Fokus auf fachkommunikative Aspekte die Relevanz diskursiver Kompetenzen als Kernkompetenz von domänenspezifischer und fachbezogener *AI Literacy* stärken. Hierzu bedarf es einer weiteren Differenzierung, evidenzbasierten Fundierung und fachlich-juristischen Ergänzung. Im Dialog von Fachkommunikator:innen aus dem Recht sowie aus weiteren Disziplinen kann dieser als Grundstein und Leitplanke für die Entwicklung von *AI (Legal) Literacy* im Fach dienen. Dieser Dialog darf als Appell an die Fachkommunikationsforschung verstanden werden, sofern sie ihr Selbstverständnis auch als Gestaltungsauftrag versteht.

Literatur

- ANDERSON, L. W. / KRATHWOHL, D. R. (2001): *A Taxonomy for Learning, Teaching and Assessing. A Revision of Bloom's Taxonomy of Educational Objectives*. New York: Longman.
- ANNAPUREDDY, R. / FORNAROLI, A. / GATICA-PEREZ, D. (2025): „Generative AI Literacy: Twelve Defining Competencies.“ *Digital Government: Research and Practice* 6.1: 1–21. DOI: 10.1145/3685680.
- ANTHROPIC (2025): *AI Fluency: Frameworks and Foundations*. <https://www.anthropic.com/ai-fluency> (29.07.2025).
- ARNTZ, M. / BAUM, M. / BRÜLL, E. / DORAU, R. / HARTWIG, M. et al. (2025): *Digitalisierung und Wandel der Beschäftigung (DiWaBe 2.0): Eine Datengrundlage für die Erforschung von Künstlicher Intelligenz und anderer Technologien in der Arbeitswelt*. Dortmund: Bundesanstalt für Arbeitsschutz und Arbeitsmedizin. DOI: 10.21934/baua:bericht20250225.
- BILLA, B. (2024): „Law as Code: Exploring Information, Communication and Power in Legal Systems.“ *CRCL* 2.1: <https://journalcrcl.org/crcl/article/view/39> (4.12.2025).
- BLOCK, M. / GÖRLICH, C. M. / LAUTENSCHLAGER, J. / MAYER, V. / STRAMM, J. et al. (2025): „KI-Kompetenzanforderungen nach Art. 4 AI Act – Juristische Analyse und

- praxisorientierte Maßnahmen.“ Whitepaper September 2025. hessian.AI. <https://zenodo.org/records/17407983> (4.12.2025).
- BLOOM, B. S. (HRSG.) (1972): *Taxonomie von Lernzielen im kognitiven Bereich*. (Beltz-Studienbuch 35). 4. Aufl. Weinheim/Basel: Beltz.
- BOWER, J. L. / CHRISTENSEN, C. M. (1995): „Disruptive Technologies. Catching the Wave.“ *Harvard Business Review*, Vol. 73: 43–53.
- BROMMER, S. / BEHRENDEN, J. / BOHLE-JUROK, U. / BUCK, I. / GIRGENSOHN, K. et al. (2023): „Wissenschaftliches Schreiben im Zeitalter von KI gemeinsam verantworten. Eine schreibwissenschaftliche Perspektive auf Implikationen für Akteur*innen an Hochschulen.“ Hrsg. Hochschulforum Digitalisierung (HFD) (Diskussionspapier Nr. 27 / November 2023). <https://hochschulforumdigitalisierung.de/news/dp-wissenschaftliches-schreiben-verantworten-ki/> (4.12.2025).
- BUNDESNETZAGENTUR (2025): „Hinweispapier KI-Kompetenzen nach Artikel 4 KI-Verordnung.“ https://www.bundesnetzagentur.de/DE/Fachthemen/Digitales/KI/7_Kompetenz/start.html (4.12.2025).
- BUSCH, A. / LUTTERMANN, K. (2024): „Einleitung: Was ist Professionskommunikation?“ *Professionskommunikation*. Hrsg. A. Busch / K. Luttermann. Hildesheim: Universitätsverlag Hildesheim. 7–13. DOI: 10.18442/fsk-5.
- BUSSE, D. (1992): *Recht als Text. Linguistische Untersuchungen zur Arbeit mit Sprache in einer gesellschaftlichen Institution*. Berlin/Boston: De Gruyter.
- BUSSE, D. (1999): „Die juristische Fachsprache als Institutionensprache (am Beispiel von Gesetzen und ihrer Auslegung).“ *Fachsprachen. Ein internationales Handbuch zur Fachsprachenforschung und Terminologiewissenschaft*. Hrsg. L. Hoffmann / H. Kalverkämper / H. E. Wiegand. Berlin/Boston: De Gruyter, 1282–1391.
- BUSSE, D. (2000): „Textsorten des Bereichs Rechtswesen und Justiz.“ *Text- und Gesprächslinguistik. Ein internationales Handbuch zeitgenössischer Forschung*. Hrsg. K. Brinker / G. Antos / W. Heinemann / S. F. Sager. Berlin: De Gruyter. 658–675.
- BUTOLLO, F. / GERBER, CH. / GÖRNEMANN, E. / GREMINGER, L. / KATZINSKI, A.-K. et al. (2024): „Die Symbiose von generativer KI und Arbeit. Erweiterung der Horizonte oder Erosion menschlicher Kompetenz?“ Hrsg. Weizenbaum-Institut e.V. Weizenbaum Discussion Paper. Berlin 41. DOI: 10.34669/WI.DP/41.
- CHAU, B. / LIVERMORE, M. (2024): „Computational Legal Studies Comes of Age.“ *EJELS* 1.1: 89–104. DOI: 10.62355/ejels.19684.

- DICKERT, TH. / OTTE, S. (2022): „Einsatz von KI und algorithmischen Systemen in der Justiz.“ Grundlagenpapier zur 74. Jahrestagung der Präsidentinnen und Präsidenten der Oberlandesgerichte, des Kammergerichts, des Bayerischen Obersten Landesgerichts und des Bundesgerichtshofs vom 23. bis 25. Mai 2022 in Rostock. Nürnberg/ Celle. https://oberlandesgericht-celle.niedersachsen.de/startseite/rechtsprechung/zivilprozess_der_zukunft_tagungsband/justiz_der_zukunft/grundlagenpapier-zum-einsatz-kunstlicher-intelligenz-in-der-justiz-215525.html (13.05.2025).
- DIETRICH, H. (2024): „Auslegen und Subsumieren mit ChatGPT. Auslegungs- und Subsumtionsfähigkeit des aktuellen Sprachmodells unter Anwendung verschiedener Prompting-Ansätze.“ *NJW (Neue Juristische Wochenschrift)* 77.29: 2092–2098.
- EFING, CH. (2024): „Registerkompetenz im Berufsfeld der Rechtsanwalts- und Notarfachangestellten.“ *Professionskommunikation*. Hrsg. A. Busch / K. Luttermann. Hildesheim: Universitätsverlag Hildesheim. 197–226. DOI: 10.18442/fsk-5.
- FELDER, E. (2008): „Das Forschungsnetzwerk ‚Sprache und Wissen‘: Zielsetzung und Inhalte.“ *Zeitschrift für Germanistische Linguistik* 36 (2): 270–276. DOI: 10.1515/ZGL.2008.020.
- FELDER, E. (2024): „Die Zeichenhaftigkeit von Gesellschaft und Wissen: Das Forschungsnetzwerk ‚Sprache und Wissen‘.“ *Sprache und Vertrauen*. Hrsg. P. Schäfer / M. Belosevic. Berlin/Boston: De Gruyter. 299–318. DOI: 10.1515/9783111452883-013.
- FELDER, E. / VOGEL, F. (2017): „Einleitung.“ *Handbuch Sprache im Recht*. Hrsg. E. Felder / F. Vogel. Berlin/Boston: De Gruyter. IX–XIX. DOI: 10.1515/9783110296198.
- FERRARI, A. (2012): „Digital Competence in Practice. An Analysis of Framework.“ *EUR 25351 EN*, Luxembourg: Publications Office of the European Union. DOI: 10.2791/82116.
- FLUCK, H.-R. (1996): *Fachsprachen. Einführung und Bibliographie*. 5., überarb. und erw. Aufl. Tübingen/Basel: UTB.
- GILSTER, P. (1997): *Digital Literacy*. New York / Chichester: Wiley.
- GÖPFERICH, S. (2008): *Textproduktion im Zeitalter der Globalisierung. Entwicklung einer Didaktik des Wissenstransfers*. 3. Aufl. Tübingen: Stauffenburg.
- HATZIUS, J. / BRIGGS, J. / KODNANI, D. / PIERDOMENICO, G. (2023): „The Potentially Large Effects of Artificial Intelligence on Economic Growth (Briggs/Kodnani).“ Hrsg. Goldman Sachs: *Global Economics Analyst. Goldman Sachs Economics Research*. https://www.ansa.it/documents/1680080409454_ert.pdf (16.12.2025).

- GRIENBERGER, K. / MATTHES, B. / PAULUS, W. (2024): „Vor allem Hochqualifizierte bekommen die Digitalisierung verstärkt zu spüren.“ *IAB-Kurzbericht. Aktuelle Analysen aus dem Institut für Arbeitsmarkt- und Berufsforschung* 5/2024. <https://doku.iab.de/kurzber/2024/kb2024-05.pdf> (25.03.2025).
- GÜLTEKIN-VÁRKONYI, G. (2025): „AI Literacy for Legal AI Systems: A Practical Approach.“ <http://arxiv.org/pdf/2505.18006v1> (4.12.2025).
- GUTIÉRREZ, J. D. (2024): „UNESCO Global Judges’ Initiative: Survey on the Use of AI Systems by Judicial Operators.“ Hrsg. UNESCO. Paris. <https://unesdoc.unesco.org/ark:/48223/pf0000389786> (4.12.2025).
- HACKL, V. / MUELLER, A. / SAILER, M. (2025): „The AI Literacy Heptagon: A Structured Approach to AI Literacy in Higher Education.“ <http://arxiv.org/pdf/2509.18900v1> (4.12.2025).
- HARTLEY, J. / JOLEVSKI, F. / MELO, V. / MOORE, B. (2025): „The Labor Market Effects of Generative Artificial Intelligence.“ DOI: [org/10.2139/ssrn.5136877](https://doi.org/10.2139/ssrn.5136877).
- HARTUNG, D. / GERLACH, L. (2025): „Legal Tech Monitor 2025. Der Markt für die Digitalisierung von Rechtsdienstleistungen in Deutschland.“ Hrsg. Legal Tech Verband e.V. <https://www.legaltechverband.de/wp-content/uploads/2025/01/Legal-Tech-Monitor-v1.0b.pdf> (4.12.2025).
- HEIDRICH, F. / SCHUBERT, K. (2020): *Fachkommunikation – gelenkt, geregelt, optimiert*. Hildesheim: Universitätsverlag Hildesheim. DOI: 10.18442/087.
- HEPP, A. / LOOSEN, W. / DREYER, ST. / JARKE, J. / KANNENGIESSER, S. et al. (2022): „Von der Mensch-Maschine-Interaktion zur kommunikativen KI.“ *Publizistik* 67.4: 449–474.
- HOFFMANN, L. (1985): *Kommunikationsmittel Fachsprache. Eine Einführung*. 2., völlig neu bearb. Aufl. Tübingen: Narr.
- IRIGHTS.LAB (2025): „Generative KI – Innovation und Recht in Arbeitsprozessen. Erkenntnisse und Befunde zu den rechtlichen und gesellschaftlichen Rahmenbedingungen beim Einsatz generativer KI.“ Working Paper. [https://www.ights-lab.de/storage/uploads/ebooks/pdfs/Gen-KI-IR-Working%20Paper%20Oktober%202025_final%20\(31.10.2025\).pdf](https://www.ights-lab.de/storage/uploads/ebooks/pdfs/Gen-KI-IR-Working%20Paper%20Oktober%202025_final%20(31.10.2025).pdf) (4.12.2025).
- JAKOBS, E.-M. (2023): „Unternehmenskommunikation in Bewegung.“ *Soziale Themen in Unternehmens- und Wirtschaftskommunikation* 35. Hrsg. Ch. M. Schmidt / S. Heinemann / V.M. Banholzer / M. Nielsen / F.U. Siems. Wiesbaden: Springer Fachmedien. 127–146. DOI: 10.1007/978-3-658-40705-6_6.

- KALVERKÄMPER, H. (1990): „Der Einfluß der Fachsprachen auf die Gemeinsprache.“ *Deutsche Gegenwartssprache. Tendenzen und Perspektiven*. Hrsg. G. Stickel. Berlin: De Gruyter. 88–133.
- KATZ, D. M. / HARTUNG, D. / GERLACH, L. / JANA, A. / BOMMARITO, M. J. (2023): „Natural Language Processing in the Legal Domain.“ *SSRN Journal*. DOI: 10.2139/ssrn.4336224.
- KNOTH, N. / DECKER, M. / LAUPICHLER, M. C. / PINSKI, M. / BUCHHOLTZ, N. et al. (2024): „Developing a Holistic AI Literacy Assessment Matrix – Bridging Generic, Domain-Specific, and Ethical Competencies.“ *Computers and Education Open* 6: 100177. DOI: 10.1016/j.caeo.2024.100177.
- KOCHER, E. (2025): „Das Schreiben von Rechtstexten als Suche nach dem Allgemeinen.“ *Recht umkämpft*. Hrsg. Forschungsgruppe „Recht – Geschlecht – Kollektivität“. Opladen/Berlin/Toronto: Barbara Budrich. 229–241. DOI: 10.3224/84743101.13
- KONERDING, K.-P. (2015): „Sprache und Wissen.“ *Handbuch Sprache und Wissen*. Hrsg. E. Felder / A. Gardt. Berlin/Boston: De Gruyter. 57–80.
- KRÜGER, J. / LISCHKA, K. (2018): „Damit Maschinen den Menschen dienen. Lösungsansätze, um algorithmische Entscheidungen in den Dienst der Gesellschaft zu stellen.“ Arbeitspapier. Hrsg. Bertelsmann Stiftung. DOI: 10.11586/2018019.
- LAUPICHLER, M. C. / ASTER, A. / SCHIRCH, J. / RAUPACH, T. (2022): „Artificial Intelligence Literacy in Higher and Adult Education. A Scoping Literature Review.“ *Computers and Education: Artificial Intelligence* 3: 100101. DOI: 10.1016/j.caeai.2022.100101.
- LAUPICHLER, M. C. / ASTER, A. / HAVERKAMP, N. / RAUPACH, T. (2023): „Development of the ‚Scale for the Assessment of Non-Experts’ AI Literacy‘ – An Exploratory Factor Analysis.“ *Computers in Human Behavior Reports* 12: 100338. DOI: 10.1016/j.chbr.2023.100338.
- LAW, N. / WOO, D. / LA TORRE, J. DE; / WONG, G. (2018): „A Global Framework of Reference on Digital Literacy Skills for Indicator 4.4.2.“ *Information Paper No. 51*. Hrsg. v. UNESCO. Montréal/Quebec. <https://unesdoc.unesco.org/ark:/48223/pf0000265403> (12.12.2025).
- LEHNEN, K. / STEINHOFF, T. (2024): „Digitales Lesen und Schreiben.“ *Handbuch Sprache und digitale Kommunikation*. Hrsg. J. Androutsopoulos / F. Vogel. Berlin/Boston: De Gruyter. 527–546. DOI: 10.1515/9783110744163-025.

- LILIENTHAL, N. / BÜCKER, S. / HERLES, CH. (2023): „Schafft Künstliche Intelligenz die Anwaltschaft ab? Einfluss von ChatGPT & Co. auf Rechtsberatung und Justiz.“ *LTO (Legal Tribune Online)*, 28.04.2023. <https://www.lto.de/recht/kanzleien-unternehmen/k/chatgpt-chatbots-kuenstliche-intelligenz-ersatz-anwaelte-richter-rechtswesen-anwaltschaft-ki>. (12.12.2025).
- LONG, D. / MAGERKO, B. (2020): „What is AI Literacy? Competencies and Design Considerations.“ *CHI '20: Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, Honolulu HI USA, April 25–30, 1–16. DOI: 10.1145/3313831.3376727.
- LOTZE, N. (2016): *Chatbots. Eine linguistische Analyse*. Bern: Lang. <https://library.oapen.org/bitstream/id/13a23aa5-f032-4642-ada3-4c50704adab6/1002758.pdf> (4.12.2025).
- MARTINETZ, S. / MARINGELE, S. (2020): „Quick Guide Legal Tech. Schritt für Schritt zur digitalen Kanzlei und Rechtsabteilung.“ Wiesbaden: Springer Gabler (Quick Guide).
- McKINSEY (2023): „The Economic Potential of Generative AI. The Next Productivity Frontier.“ <https://www.mckinsey.de/~media/mckinsey/locations/europe%20and%20middle%20east/deutschland/news/presse/2023/2023-06-14%20mgi%20genai%20report%2023/the-economic-potential-of-generative-ai-the-next-productivity-frontier-vf.pdf> (12.12.2025).
- McKINSEY & COMPANY (2025): „What is an AI agent?“ <https://www.mckinsey.com/featured-insights/mckinsey-explainers/what-is-an-ai-agent> (12.12.2025).
- MEIER, CH. (2025): „Generative KI, Wissensarbeit & Personalentwicklung. Aktuelle Studienergebnisse und Handlungsfelder für Verantwortliche in ‚Learning & Development‘ / Personalentwicklung.“ *SCIL Arbeitsbericht 34*. Universität St. Gallen. https://ibb.unisg.ch/fileadmin/user_upload/HSG_ROOT/Institut_IBB/Kompetenzzentrum_SCIL/Publikationen_2025-2028/SCIL-AB-34_GenKI-WissensArbeit-PE_Vs2025-04-18.pdf (17.04.2025).
- MEISENBACHER, ST. / MACHNER, N. / VLADIKA, J. / MATTHES, F. (2024): *Legal AI Use Case Radar 2024 Report*. Technical University of Munich. https://legal-ai-radar.de/pdfs/Legal_AI_Use_Case_Radar_Report_2024.pdf (4.12.2025).
- MIELKE, B. (2025): „Künstliche Intelligenz in der Justiz – ein Update. Einblicke in die derzeit wichtigsten Einsatzfelder und Projekte.“ *Legal-Tech.de magazin 2*: 19–27. https://www.legal-tech.de/Magazine/FFI_Legal_Tech-Magazin_02-25.pdf (12.12.2025).

- MÜLLER, F. (1994): *Strukturierende Rechtslehre*. 2., vollst. neu bearb., erg. Aufl. Berlin: Duncker & Humblot.
- MUNDORF, M. / SEIDL, E. (2025): „Generative KI in Fachkommunikation und Translation: Herausforderungen für Aus- und Weiterbildung in Wirtschaft und Recht.“ *New Generation Communication: Die Kommunikation in einer veränderten Welt*. Hrsg. Vidya Allemann-Ravi. Wiesbaden: Springer Fachmedien Wiesbaden. 277–305. DOI: 10.1007/978-3-658-47591-8_10.
- NG, D. T. K. / LEUNG, J. K. L. / CHU, S. K. W. / QIAO, M. S. (2021): „Conceptualizing AI Literacy: An Exploratory Review.“ *Computers and Education: Artificial Intelligence* 2: 100041. DOI: 10.1016/j.caeai.2021.100041.
- OECD (2019): *OECD Future of Education and Skills 2030. OECD Learning Compass 2030. A Series of Concept Notes*. Paris: OECD Publishing. https://www.oecd.org/content/dam/oecd/en/about/projects/edu/education-2040/1-1-learning-compass/OECD_Learning_Compass_2030_Concept_Note_Series.pdf (12.12.2025).
- OECD (2025): *Empowering Learners for the Age of AI. An AI Literacy Framework for Primary and Secondary Education*. Review draft (May 2025). Paris: OECD Publishing. <https://ailiteracyframework.org/> (12.12.2025).
- RHEIN, L. (2015): *Selbstdarstellung in der Wissenschaft. Eine linguistische Untersuchung zum Diskussionsverhalten von Wissenschaftlern in interdisziplinären Kontexten*. Wissen – Kompetenz – Text 8). Frankfurt a. M. / Bern: Lang.
- ROELCKE, TH. (2020): *Fachsprachen*. (ESV basics37). Berlin: Erich Schmidt.
- ROELCKE, TH. (2023): „Funktionen fachlicher Kommunikation.“ *Glottology* 14.1: 111–124. DOI: 10.1515/glot-2023-2005.
- ROELCKE, TH. / BREEZE, R. / ENGBERG, J. (Hrsg.) (2025): *Specialized Communication. An International Handbook*. (HSK 47.1). Berlin/Boston: De Gruyter. DOI: 10.1515/9783110672633.
- SACK, D. / KRAYER, L. / WILES, E. / ABBADI, M. / AWASTHI, U. / KENNEDY, R. / ARNOLDS, C. / CANDELON, F. (2024): „GenAI Doesn’t Just Increase Productivity. It Expands Capabilities.“ <https://www.bcg.com/publications/2024/gen-ai-increases-productivity-and-expands-capabilities>. (12.12.2025).
- SCHNEIDER, J. G. (2024): „Intelligible Texturen. Welche Rolle kann ChatGPT bei der Aufsatzbewertung spielen?“ Hrsg. v. Virtuelles Kompetenzzentrum – Schreiben lehren und lernen mit Künstlicher Intelligenz (VK:KIWA) (VKI:KIWA). <https://zenodo.org/records/10877034>.

- SCHUBERT, K. (2003): „Integrative Fachkommunikation.“ Übersetzen und Dolmetschen. Modelle, Methoden, Technologie. Hrsg. K. Schubert. Tübingen: Narr. 225–256.
- SCHUBERT, K. (2007): „Wissen, Sprache, Medium, Arbeit. Ein integratives Modell der ein- und mehrsprachigen Fachkommunikation.“ Tübingen: Narr.
- SCHULZ, W. / SCHMEES, J. (2022): „Möglichkeiten und Grenzen der Künstlichen Intelligenz in der Rechtsanwendung.“ *Wissen und Recht*. Hrsg. I. Augsberg / G. F. Schuppert. Baden-Baden: Nomos. 561–593.
- SEUFERT, S. / MEIER, CH. (2024): „Hybride Intelligenz: Zusammenwirken von menschlicher und maschineller Intelligenz.“ *Generative Künstliche Intelligenz. ChatGPT und Co für Bildung, Wirtschaft und Gesellschaft*. Hrsg. S. Seufert / S. Handschuh. Freiburg: Schäffer-Poeschel. 57–72.
- STEINHOFF, T. (2023): „Der Computer schreibt (mit): Digitales Schreiben mit Word, Whatsapp, ChatGPT & Co. als Koaktivität von Mensch und Maschine.“ *MIDU* 5.1: 1–16. DOI: 10.18716/ojs/midu/2023.1.4.
- SUCHOWOLEC, K. (2024): „Fachkommunikative Modelle.“ *Handbuch Deutsch als Fach- und Fremdsprache*. Hrsg. M. Szurawitzki / P. Wolf-Farré. Berlin/Boston: De Gruyter. 39–54. DOI: 10.1515/9783110690279-004.
- TOMLINSON, K. / JAFFE, S. / WANG, W. / COUNTS, S. / SURI, S. (2025): „Working with AI: Measuring the Occupational Implications of Generative AI.“ *arXiv*. DOI: 10.48550/arXiv.2507.07935.
- TU, P. (2024): „Die Wellen des Legal Tech. Legal Tech – Die Rechtsbranche im Wandel.“ *Legal Tech Verzeichnis*. <https://legal-tech-verzeichnis.de/fachartikel/die-wellen-des-legal-tech/> (12.12.2025).
- VOGEL, F. (2024): „Der Richter, (k)ein Bot?!“ *Digitalisierung des Rechts*. Hrsg. L. Wörner / R. Wilhelmi / J. Glöckner / M. Breuer / S. Behrendt. Berlin/Boston: De Gruyter. 9–26. DOI: 10.1515/9783111343341-004
- WEINERT, FRANZ E. (2002): „Vergleichende Leistungsmessung in Schulen – eine umstrittene Selbstverständlichkeit.“ *Leistungsmessungen in Schulen*. Hrsg. F. E. Weinert. Weinheim/Basel: Beltz. 17–31.
- WILDER, N. / WESSELS, D. / KLEIN, A. / MUNDORF, M. / GRÖPLER, J. (2021): *Who is Leading Whom in Research in the Age of Artificial Intelligence?* (Translated and adapted by Katharina Miller). <https://leadingwithai.substack.com/p/>

- who-is-leading-whom-in-research-in?r=t32sk&utm_campaign=post&utm_medium=email&utm_source= (14.10.2021).
- WIMMER, R. (2008): „Weltansichten aus sprachlicher und rechtlicher Perspektive. Zur Ontisierung von Konzepten des Rechts.“ *Verständlichkeit als Bürgerrecht? Die Rechts- und Verwaltungssprache in der öffentlichen Diskussion*. Hrsg. K. M. Eichhoff-Cyrus. Mannheim: Dudenverlag. 81–95.
- WIMMER, R. (2009): „Zur Verflechtung von Spracharbeit und Rechtsarbeit in der EU.“ *Muttersprache* 3: 234–239.
- WÖRNER, L. / WILHELMI, R. / GLÖCKNER, J. / BREUER, M. (2024): „Digitalisierung des Rechts. Eine Einführung.“ *Digitalisierung des Rechts*. Hrsg. L. Wörner / R. Wilhelmi / J. Glöckner / M. Breuer / S. Behrendt. Berlin/Boston: De Gruyter. 1–7. DOI: 10.1515/9783111343341-003.
- WORLD ECONOMIC FORUM (2025): *Future of Jobs Report 2025. Insight Report January 2025*. Geneva. <https://www.weforum.org/publications/the-future-of-jobs-report-2025/digest/> (12.12.2025).
- ZERFAB, A. / BUHMANN, A. / TENCH, R. / VERČIČ, D. / MORENO, Á. (2021): *European Communication Monitor 2021. CommTech and Digital Infrastructure, Video-Conferencing, and Future Roles for Communication Professionals. Results of a survey in 46 countries*. EUPRERA/EACD. Brüssel.

Auswirkungen von KI auf Terminologearbeit im institutionellen Kontext

Workflows, Rollen und Kompetenzen

Abstract

Due to the rapid advancement of Artificial Intelligence (AI) and its increasingly easy accessibility – even for people without programming skills – AI-based tools are being used in more and more specialist areas. The field of communication for specific purposes is also affected by this development. This article examines the impact of AI on practical terminology work, particularly in an institutional setting. The exploratory study is based on qualitative expert interviews and analyses which steps in terminology work can be supported by AI and to what extent AI is already integrated into the terminology workflow. It also examines how the use of such technologies affects the roles and competence requirements of terminology experts.

Keywords KI, Kompetenzen, Rollen, Terminologearbeit, Workflows

1 Einleitung

Durch die rasante Entwicklung der Künstlichen Intelligenz (KI) und ihren einfachen Zugang, selbst für Menschen ohne Programmierkenntnisse, werden immer mehr KI-basierte Tools von Fachleuten in verschiedenen Bereichen eingesetzt, so auch von Sprachexpert*innen wie Übersetzer*innen oder Termino-

log*innen. Dieser Beitrag untersucht, wie KI die praktische Terminologiearbeit in institutionellen Einrichtungen verändert. Terminologiearbeit ist im Kontext der Fachkommunikation und Fachübersetzung sehr wichtig, da mit Hilfe der Terminologiearbeit Ressourcen wie mehrsprachige Terminologiewörterbücher, Terminologiedatenbanken und Glossare erstellt und aktualisiert werden, die dann beispielsweise von Fachexpert*innen und Übersetzer*innen konsultiert werden. Einige Autor*innen haben bereits Experimente durchgeführt, welche Arbeitsschritte im Terminologieworkflow durch KI-basierte Tools automatisiert oder unterstützt werden können: Definitionen schreiben (Reineke 2023; San Martín 2024), Erstellung von Konzeptsystemen (Łukasik 2023; Reineke 2023), Termextraktion (Łukasik 2023; Varga 2024) und die Ermittlung von Synonymen (Reineke 2023). Es gibt jedoch keine aktuellen Über-sichten darüber, wie KI-basierte Tools in der praktischen Terminologiearbeit in Institutionen tatsächlich eingesetzt werden und welche Auswirkungen der Einsatz von KI auf die Rollen, Aufgaben, Workflows und Kompetenzen aller an der Terminologiearbeit beteiligten Personen hat. Um diese Lücke zu schließen, wurde die vorliegende interviewbasierte explorative Studie durchgeführt und die folgenden Forschungsfragen genauer beleuchtet: Welche Arbeitsschritte im Terminologieworkflow können bereits durch KI unterstützt oder automa-tisiert werden und wie verändern sich Rollenprofile und Kompetenzen von Terminologieexpert*innen?

Der vorliegende Artikel gliedert sich in eine Einleitung, danach wird unter Punkt 2 auf die Arbeitsschritte in der Terminologiearbeit eingegangen. Unter Punkt 3 wird der Forschungsstand zum Einsatz von KI in der Terminologiearbeit diskutiert und danach werden die Rollenprofile und die Kompetenzen im Zeitalter von KI beleuchtet. Unter Punkt 5 werden die Methode der vorliegenden Studie beschrieben, die Ergebnisse werden unter Punkt 6 präsentiert und diskutiert. Der Artikel schließt mit einem Fazit und Ausblick.

2 Terminologieworkflows

In Institutionen, wie z. B. in der öffentlichen Verwaltung, in Universitäten, in internationalen Organisationen, kann Terminologearbeit für unterschiedliche Zwecke durchgeführt werden, etwa zur Verbesserung der Kommunikation innerhalb eines Fachbereichs oder einer Nutzungsgemeinschaft oder zur Verbesserung und Qualitätssicherung von Übersetzungen. Ganz gleich zu welchem Zweck Terminologearbeit durchgeführt wird, im Workflow können spezielle Arbeitsschritte unterschieden werden (Chiocchetti et al. 2023):

Der erste Schritt ist die Bedarfsanalyse, bei der erhoben wird, welche Art von Terminologearbeit durchgeführt werden soll. Falls noch kein Terminologieverwaltungssystem bestehen sollte, würde als nächster Schritt das Design und die Implementierung eines solchen Systems folgen. Danach erfolgt die Dokumentation, bei der die textuellen Ressourcen gesammelt werden, die dann im nächsten Schritt für die Termextraktion verwendet werden.

Bei der Termextraktion geht es um die Identifizierung möglicher Termini aus den zuvor gesammelten textuellen Ressourcen, die dann in die terminologische Ressource aufgenommen werden sollen. Das Herausfiltern kann entweder manuell erfolgen, indem das Material gelesen wird, oder toolgestützt. Wenn die Termextraktion toolgestützt erfolgt, muss die Termkandidatenliste manuell bereinigt werden, bevor Termeinträge erstellt werden können. Dabei werden die Termini mit zusätzlichen Informationen, entsprechend dem Datenmodell, angereichert, z. B. mit Definitionen, Quellenangaben, Kontexten, Gebrauchsinformationen, Synonymen und Varianten, Äquivalenten in anderen Sprachen.

Beim Schritt der Revision und Validierung werden die Einträge formal, linguistisch und inhaltlich geprüft. Die formale Überprüfung übernehmen meist Terminologiemanager*innen oder Terminolog*innen, die kontrollieren, ob die Einträge vollständig sind, ob alle Informationen den richtigen Datenkategorien zugeordnet sind etc. Bei der linguistischen Revision, die meist von Muttersprachler*innen durchgeführt wird, stehen sowohl die Rechtschreibung als auch neugebildete Termini im Blickpunkt der Kontrolle.

Bei der inhaltlichen Revision, meist durch Fachexpert*innen ausgeführt, werden die Definitionen, Synonyme und Äquivalente überprüft. Nach diesem

Schritt können die Daten veröffentlicht und verbreitet werden. Oft beinhaltet dies auch die Konvertierung der Daten, um sie in unterschiedlichen Formaten bereitzustellen. Weiters fällt auch das Bereitstellen der Daten in Datenrepositorien unter den Veröffentlichungsschritt (Wissik 2024; Wissik 2025).

Nach der Veröffentlichung der Daten ist die Terminologearbeit aber noch nicht abgeschlossen, denn damit die Inhalte weiterhin relevant und nutzbar bleiben, müssen sie permanent gepflegt und aktualisiert werden: z. B. durch die Bereinigung von Dubletten, durch das Zusammenführen von einzelnen Einträgen, durch das Aktualisieren von Definitionen. Auch auf der technischen Ebene müssen die Terminologieverwaltungsprogramme gewartet werden (Chiocchetti et al. 2023).

3 KI und Terminologearbeit

Der Einsatz von computergestützten Methoden in der Terminologearbeit ist nicht neu und reicht in die frühen 1990er Jahre zurück, als die ersten Terminologieverwaltungsprogramme zum Einsatz kamen. Immer mehr Tools werden für die unterschiedlichen Arbeitsschritte eingesetzt, um diese zu automatisieren, z. B. Tools für die Extraktion von Termini, Tools für die Extraktion von Definitionen oder auch maschinelle Übersetzung. Bisher waren die Tools sehr spezifisch für einen bestimmten Arbeitsschritt und basierten auf unterschiedlichen Methoden (z. B. regelbasiert, statistisch oder basierend auf neuronalen Netzwerken). KI-basierte Tools können die Terminologearbeit transformieren, indem sie mühsame und zeitaufwändige Prozesse (z. B. Termextraktion oder das Verfassen von Definitionen) automatisieren. Insbesondere durch den Einsatz von generativen KI-basierten Tools wie ChatGPT oder die Integration dieser in bestehende Terminologieverwaltungsprogramme können viele unterschiedliche Schritte mit ein und demselben Tool unterstützt werden.

Seit der Markteinführung von ChatGPT im November 2022 hat das Thema auch im Bereich der praktischen Terminologearbeit Beachtung gefunden und zahlreiche Experimente mit generativer KI, insbesondere mit unterschiedlichen Versionen von ChatGPT, wurden durchgeführt und deren Ergebnisse

publiziert. Der Einsatz von KI in diesen Experimenten reicht von der Identifizierung von Dokumenten zur Integration in Korpora (Nagendra/Chandra 2022) über Termextraktion (Reineke 2023) und Extraktion von fachsprachlichen Kollokationen (Giacomini 2025) bis hin zum Verfassen von Definitionen (San Martín (2024), der Identifizierung von Äquivalenten (Massion 2024) oder der Übersetzung von Neologismen (Lerner/Yvon 2025). Reineke (2023) untersucht mit ChatGPT für das Deutsche die Erstellung von Definitionen, den Aufbau von Begriffssystemen und die Bestimmung von Synonymen. Laut Reineke (2023), scheinen die Resultate auf den ersten Blick nützlich zu sein, auf den zweiten Blick sind die Ergebnisse aber oft inkonsistent: Dieselbe Anfrage ergibt unterschiedliche Resultate und sie enthalten oft Halluzinationen, d. h. falsche Ergebnisse. Łukasik (2023) führt ähnliche Tests auf Englisch und Polnisch durch, mit besseren Ergebnissen für das Englische. Sánchez-Gijón/Palenzuela-Badiola (2023) stellen fest, dass GPT-4 bei Terminologieextraktion und Äquivalenten bessere Ergebnisse liefert als GPT-3 und Varga (2024) experimentiert mit dem Sprachenpaar Spanisch–Rumänisch, wobei ChatGPT bessere Ergebnisse für das Spanische liefert. Giacomini (2024) experimentiert mit ChatGPT und Perplexity, um nicht hierarchische Begriffsbeziehungen zu identifizieren. Heinisch (2025b) zeigt anhand systemgebundener Domänen und Sprachvarietäten die Grenzen von Large Language Models (LLMs) auf, auf denen Anwendungen wie ChatGPT basieren. Für sehr spezifische und systemgebundene Fachbereiche in Kombination mit spezifischen Sprachvarietäten verfügen die LLMs oft nicht über ausreichende Trainingsdaten aus dem spezifischen Bereich und der spezifischen Varietät und daher liefern sie fehlerhafte Ergebnisse. Heinisch (2025) experimentiert mit Termini aus dem Bereich des Hochschulwesens für das Sprachenpaar Englisch und Österreichische Varietät des Deutschen und hier werden von ChatGPT z. B. Bundesdeutsche Termini vorgeschlagen anstelle der in Österreich verwendeten Termini.

Bei den meisten Experimenten wurden generische LLMs wie ChatGPT oder Perplexity verwendet, die zwar sehr vielseitig sind, weil sie auf unterschiedlichsten Daten trainiert wurden, aber in keinem Bereich auf sehr fachspezifische Trainingsdaten zurückgreifen können. Dies kann aber bei bestimmten terminologischen Fragestellungen zu sehr allgemeinen oder aber auch falschen Ergebnissen führen (Heinisch 2025b). Daher werden für einige Fachbereiche

auch fachspezifische LLMs erstellt, wobei generische LLMs durch Feinabstimmung mit fachspezifischen Daten optimiert werden und daher fachspezifischere Antworten geben können. Es gibt z. B. fachspezifische LLMs für den Bereich Medizin (Med Palm 2 [Singhal et al. 2025]) oder für den Bereich Klima (ClimateBert [Webersinke et al. 2021]). Da sich das Fachwissen sehr schnell ändern kann und auch bei fachspezifischen LLMs Trainingsdaten bis zu einem bestimmten Datum verwendet werden, können auch fachspezifische LLMs Wissenslücken aufweisen. Nicht für alle Fachbereiche gibt es fachspezifische LLMs, und da das Feinabstimmen von LLMs oft sehr ressourcenintensiv ist, ist es sinnvoll, sich auch andere Methoden anzuschauen, bei denen fachspezifische Informationen und LLMs miteinander verbunden werden können. Eine solche Möglichkeit bietet die sogenannte *Retrieval-Augmented-Generation* (RAG)-Technologie, bei der die Ausgabe von LLMs durch die Integration einer externen Wissensbasis, die nicht Teil der Trainingsdaten war, optimiert wird (Lewis et al. 2020). Als externe Wissensbasis können Datenbanken, Textkorpora, Ontologien oder auch Terminologiedatenbanken dienen. Diese Technologie hat im Gegensatz zu fachspezifischen LLMs den Vorteil, dass immer laufend aktualisierte Daten integriert werden können. Khemakhem et al. (2025) diskutieren unterschiedliche Ansätze, wie terminologische Ressourcen in KI-Systeme eingebunden werden können, um die Ergebnisse zu verbessern. Wenn Terminologiedatenbanken integriert werden, spricht man von *Terminology Augmented Generation* (TAG) (Fleischmann/Lang 2025).

Bei den meisten Experimenten wurden die üblichen Chatinterfaces genutzt mit ihren bekannten Einschränkungen (z. B. Zeichenbeschränkung bei Prompts, Größenbeschränkung beim Hochladen von Dateien), daher betonen mehrere Autor*innen (u. a. Łukasik 2023; Reineke 2023; San Martín 2024) die Notwendigkeit, KI stärker in bestehende Terminologiewerkzeuge zu integrieren oder spezialisierte Tools zu entwickeln, die für die praktische Terminologearbeit angepasst sind.

4 Rollenprofil und Kompetenzen im Zeitalter von KI

In der Terminologiewerkarbeit gibt es unterschiedliche Beteiligte in unterschiedlichen Rollen, darunter Terminolog*innen, Terminologiemanager*innen, Revisor*innen, Fachexperten, IT-Experten und Nutzer*innen (siehe TerminOrgs 2016; RaDT 2020; Chiochetti et al. 2023). Alle Beteiligten haben unterschiedliche Aufgaben und demzufolge auch unterschiedliche Kompetenzen. Bevor KI allgegenwärtig wurde, waren für Terminolog*innen folgende Kompetenzen erforderlich (RaDT 2020): Kenntnisse der Terminologiewissenschaft, Kenntnisse der Methoden der praktischen Terminologiewerkarbeit, Recherchekompetenz und kritische Quellenbewertung, technologische Kompetenz (effektive Nutzung von für die Terminologiewerkarbeit notwendigen Tools, z. B. Terminologieverwaltungsprogrammen, Termextraktionstools), Sachkompetenz (Fachwissen in den Fachgebieten, für die Terminologiewerkarbeit betrieben wird), Sprachkompetenz in den Sprachen, in denen gearbeitet wird. Weiters war auch noch die Datenkompetenz relevant, vor allem das Wissen um unterschiedliche Datenformate und Standards, sowie die Fähigkeit, Daten von einem Format in ein anderes zu konvertieren.

Seit den rapiden Entwicklungen im Bereich der generativen KI befassen sich unterschiedliche Publikationen mit der Veränderung der Rollenprofile der Beteiligten im Terminologieworkflow, wobei sich die meisten auf das Profil von Terminolog*innen konzentrieren (Heinisch 2025a; Heyn/Allard 2025). Wissik (2025c) diskutiert auch die Entwicklung einer neuen Rolle, nämlich die der Technologiemanager*innen. Außerdem diskutieren Chiochetti et al. (2023) Maschinen als Beteiligte im Terminologieworkflow.

Da sich durch den Einsatz von KI die Rollen und die Aufgaben verändern, müssen auch die Kompetenzen angepasst werden. Als Beispiel soll hier eine Umfrage zum Thema Nutzung von KI unter Forschenden und Lehrenden in der CLARIN Community erwähnt werden, bei der die Teilnehmenden auch nach den neuen Kompetenzen gefragt wurden, die sie für die Nutzung von KI benötigen. Die Mehrheit der Befragten gab an, dass sie ihre technischen Kompetenzen vertiefen müssten, im Speziellen die Prompt Engineering Kompetenz, das allgemeine Wissen darüber, wie KI-Modelle funktionieren, und

deren Vor- und Nachteile, aber auch Datenmanagementkompetenzen wurden genannt (van der Lek et al. 2025). Diese Umfrageergebnisse schneiden bereits zwei wichtige Kompetenzen im Zeitalter von KI an: zum einen die Datenkompetenz („data literacy“) und zum anderen die KI-Kompetenz („AI literacy“).

Bei der Datenkompetenz geht es um die „Fähigkeiten, Daten auf kritische Art und Weise zu sammeln, zu managen, zu bewerten und anzuwenden“ (Schüller et al. 2019). Die Datenkompetenz wird auch als Grundlage für die KI-Kompetenz gesehen, da die Vorbereitung von KI-fähigen Daten, d. h. von qualitativ hochwertigen Daten, die so aufbereitet sind, dass sie von KI-Systemen direkt verwendet werden können, in erster Linie als Teil der Datenkompetenz zählt. Damit terminologische Ressourcen KI-fähig sind, müssen die terminologischen Daten so aufbereitet und strukturiert sein, dass sie effektiv von KI verarbeitet, analysiert und genutzt werden können. Dazu gehört auch, dass z. B. die Daten in einem bestimmten Format vorliegen. Obwohl LLMs mit terminologischen Datenformaten wie TBX (TermBase eXchange) umgehen können, sind Formate wie JSON laut Fleischmann/Lang (2025) besser geeignet. Es kann jedoch zwischen den einzelnen LLMs Unterschiede geben, welche Formate verwendet werden sollten. Die Datenkompetenz von Terminologiemanager*innen, Technologiemanager*innen oder Terminolog*innen beinhaltet demzufolge auch die Kompetenz, KI-fähige terminologische Daten bereitzustellen und im Zuge dessen die terminologischen Daten von einem Format in ein anderes Format überführen zu können.

Während die Datenkompetenz die grundlegende Voraussetzung für eine effektive KI-Nutzung bildet, umfasst die KI-Kompetenz einen breiteren Kontext hinsichtlich der Funktionsweise und des Potenzials von KI. KI-Kompetenz beinhaltet die Fähigkeiten, KI-Technologien kritisch zu evaluieren, effektiv mit KI zu kommunizieren und zusammenzuarbeiten als auch KI als Werkzeug online, zu Hause und am Arbeitsplatz nutzen zu können (vgl. Long/Magerko 2020). In Strauß (2024) finden sich Kompetenzen für Wissensarbeiter*innen, zu denen auch die Terminologieexpert*innen zählen. Heinisch (2025a) diskutiert KI-Kompetenzen, sie nennt sie LLM Literacy für Terminolog*innen. Die Kompetenzen bauen auf dem vierstufigen Kompetenzmodell von Hibbert et al. (2024) auf: Die erste Stufe umfasst das Wissen über die Funktionsweise von KI-Modellen, die zweite Stufe betrifft die Anwendung von KI-Modellen, die

dritte Stufe umfasst die Analyse und Evaluierung der KI-Modelle und die vierte Stufe das Erstellen und Anpassen von KI-Modellen. Zu den KI-Kompetenzen von Terminolog*innen zählen laut Heinisch (2025a): Modellverständnis sowie Wissen um die Stärken und Schwächen von LLMs, Wissen um die Grenzen der Sprachunterstützung, Kompetenz zur Validierung der Qualität, Aktualität und Validität von KI-Outputs, Abschätzen der Anwendungsmöglichkeiten im Workflow (wann ist es sinnvoll KI einzusetzen, wann nicht) sowie Wissen über ethische und datenschutzrechtliche Aspekte. Auch Prompt Engineering zählt zu den KI-Kompetenzen im Bereich der Anwendung von KI-Modellen. Über Prompts kommuniziert und interagiert der Mensch mit dem KI-Modell. Es sind Anweisungen, die an ein KI-Modell gerichtet werden, um ein spezielles Ergebnis zu erzielen. Prompts können in Form von Texten aber auch in Form von Bildern, Audiodaten oder Codes erfolgen. Die Qualität und die Präzision eines Prompts sind entscheidend für die Genauigkeit und die Relevanz der KI-Ausgabe (Long et al. 2025).

5 Methode

Die hier beschriebene Studie ist Teil des Projekts über Terminologieworkflows im 21. Jahrhundert und besteht aus zwei Teilen: Zuerst wurden 15 qualitative Online-Fachexpert*inneninterviews (Meuser/Nagel 1991) mit in Institutionen tätigen Terminologieexpert*innen im Jahr 2023 via Zoom durchgeführt und als Follow-up wurden dieselben Expert*innen zu E-Mail-Interviews (Ratislavová/Ratislav 2014) im Jahr 2025 eingeladen. Dieser Einladung folgten 13 Expert*innen. Die genauen Profile der Interviewteilnehmer*innen sind in Tabelle 1 aufgelistet. Terminologearbeit in Unternehmen ist nicht Gegenstand dieser Untersuchung. Die ausführliche Beschreibung der Methode findet sich in Wissik (2025a) und Wissik (2025c), da einzelne Aspekte des Projekts schon publiziert wurden.

Die leitfadengestützten Online-Zoom-Interviews waren umfangreich und umfassten eine Vielzahl von Themen¹, beginnend mit allgemeinen Fragen zur Institution und dazu, welche Art der Terminologiarbeit durchgeführt wird, sowie zur Rolle der Befragten über Fragen zum Einsatz von Technologie bis zum Workflow. Die Fragen² in den E-Mail-Interviews fokussierten sich auf die Nutzung von generativen KI-Anwendungen, ob mit generativen KI-Anwendungen bereits experimentiert wurde und wenn ja, für welche Arbeitsschritte und wie die Erfahrungen waren und welche Rolle generative KI in der täglichen Terminologiarbeit spielt (Wissik 2025c).

Tabelle 1: Übersicht Interviewteilnehmer*innen

ID	Rolle	Institution	Art der durchgeführten Terminologiarbeit in der Institution	Teil I	Teil II
INT 1	IT-Experte im Bereich Terminologiesoftware	Wissenschaft/ Forschung	systematische Terminologiarbeit / punktuelle Terminologiarbeit / Vorarbeit für die Normierung	Ja	Ja
INT 2	Technologie-manager*in	Wissenschaft/ Forschung	systematische Terminologiarbeit / punktuelle Terminologiarbeit / Vorarbeit für die Normierung	Ja	Ja

.....

- 1 Der Leitfaden zu den Online-Zoom-Interviews ist auf Zenodo unter <https://doi.org/10.5281/zenodo.11144968> abrufbar.
- 2 Die Fragen für die E-Mail-Interviews sind auf Zenodo unter <https://doi.org/10.5281/zenodo.17762721> abrufbar.

ID	Rolle	Institution	Art der durchgeführten Terminologiearbeit in der Institution	Teil I	Teil II
INT 3	Leitung Terminologieabteilung/ Terminolog*in	Wissenschaft/ Forschung	systematische Terminologiearbeit / punktuelle Terminologiearbeit	Ja	Ja
INT 4	Mitglied einer Terminologie-Kommission	Wissenschaft/ Forschung	systematische Terminologiearbeit / punktuelle Terminologiearbeit	Ja	Nein
INT 5	Leitung Terminologie und Rechtsübersetzungsabteilung, Stellvertretende Leitung Entwicklung	Verwaltung (Staatsebene)	übersetzungsorientierte Terminologiearbeit / a posteriori Terminologiearbeit / Vorarbeit für die Normierung / Koordination der Terminologie zwischen EU und Mitgliedsstaat(en)	Ja	Ja
INT 6	Terminolog*in	Verwaltung (Regionalebene)	Vorarbeit für die Normierung / punktuelle Terminologiearbeit / systematische Terminologiearbeit	Ja	Ja
INT 7	Terminolog*in	Wissenschaft/ Forschung	systematische Terminologiearbeit / punktuelle Terminologiearbeit	Ja	Ja
INT 8	Leitung Projektmanagement/ Terminolog*in	Verwaltung (Regionalebene)	systematische Terminologiearbeit / punktuelle Terminologiearbeit	Ja	Ja

ID	Rolle	Institution	Art der durchgeführten Terminologearbeit in der Institution	Teil I	Teil II
INT 9	Terminolog*in	Wissenschaft/ Forschung	systematische Terminologearbeit / punktuelle Terminologearbeit / Vorarbeit für die Normierung	Ja	Ja
INT 10	Terminologiekoodinator*in/ Terminologiemanager*in	Europäische Institution	übersetzungsorientierte Terminologearbeit / textbasierte Terminologearbeit	Ja	Ja
INT 11	Leitung Terminologieabteilung/ Terminolog*in	Internationale Organisation	übersetzungsorientierte Terminologearbeit / systematische Terminologearbeit (projektbasiert) / punktuelle Terminologearbeit	Ja	Ja
INT 12	Technologie-manager*in	Europäische Institution	übersetzungsorientierte Terminologearbeit / textbasierte Terminologearbeit	Ja	Ja
INT 13	Leitung Terminologieabteilung/ Terminolog*in	Internationale Organisation	übersetzungsorientierte Terminologearbeit / systematische Terminologearbeit (projektbasiert)	Ja	Nein

ID	Rolle	Institution	Art der durchgeführten Terminologearbeit in der Institution	Teil I	Teil II
INT 14	Terminologiekordinator*in/ Terminologiemanager*in	Wissenschaft/ Forschung	systematische Terminologearbeit	Ja	Ja
INT 15	Leitung Terminologieabteilung/ Terminolog*in	Verwaltung (Regional- ebene)	textbasierte Terminologearbeit / a posteriori Terminologearbeit / manchmal proaktive Terminologearbeit / punktueller Terminologearbeit	Ja	Ja

Die Zoom-Interviews mussten transkribiert werden, während die E-Mail-Interviews schon in schriftlicher Form vorlagen. Das Interviewkorpus umfasst mehr als 105.000 Tokens (Interviewkorpus I 100.497 Tokens und Interviewkorpus II 5.066 Tokens) und wurde mit Hilfe der qualitativen Inhaltsanalyse (Kuckartz 2014) untersucht. Die zu annotierenden Datenkategorien wurden sowohl deduktiv als auch induktiv gebildet.

Wie beschrieben, ist die Stichprobe eher klein und kann als Einschränkung der Studie angesehen werden. Die Stichprobe deckt jedoch die gängigsten Ansätze der Terminologearbeit und die gängigsten Szenarien ab, in denen Terminologearbeit in Institutionen durchgeführt wird (siehe Tabelle 1). Die Studie konzentriert sich vorrangig auf Institutionen in Europa, es sind aber auch zwei internationale Institutionen in der Stichprobe enthalten. Darüber hinaus bedeuten mehr Interviews nicht unbedingt mehr Daten, wie die Experimente von Guest et al. (2006) zeigen, wo bereits 12 qualitative Interviews ausreichen, um eine Datensättigung zu erreichen.

6 Ergebnisse

6.1 KI-gestützte Arbeitsschritte im Terminologieworkflow

Fast alle Interviewteilnehmer*innen (außer INT 14) haben generative KI-Anwendungen wie ChatGPT für die Terminologearbeit getestet. In allen teilnehmenden Institutionen war zum Zeitpunkt der Interviews die Verwendung von KI-Anwendungen nicht obligatorisch, sondern es lag im Ermessen der Terminologieexpert*innen, ob und für welche Arbeitsschritte KI-Anwendungen eingesetzt werden (Wissik 2025c). Die folgenden KI-Anwendungen wurden in den Interviews erwähnt: ChatGPT, MS Copilot, DeepSeek und DeepL, ein KI unterstütztes maschinelles Übersetzungssystem. Zwei Interviewteilnehmer*innen gaben an, die öffentlichen generativen Modelle für die Arbeit nicht nutzen zu dürfen und nur auf die Inhouse-Systeme zurückgreifen zu können. Wenn wir die Ergebnisse aus den Interviews nach den zuvor beschriebenen Workflowschritten ordnen, ergibt sich das folgende Bild (siehe auch Wissik 2025c): Für die Arbeitsschritte der Bedarfsanalyse, des Designs und der Implementierung eines Terminologieverwaltungssystems wurden KI-Anwendungen im Interviewkorpus nicht erwähnt. Zwei Interviewteilnehmer*innen, die in die Entwicklung von Terminologieverwaltungssystemen involviert sind, gaben an, KI-Anwendungen in naher Zukunft in das Terminologieverwaltungssystem zu integrieren (INT 1, INT 2), jedoch beim Programmieren selber keine KI-Anwendungen zu nutzen.

Der erste Schritt, bei dem KI laut dem Interviewkorpus eine Rolle spielt, ist der Schritt der Dokumentation. KI wird verwendet, um fachsprachliche Quellen zu suchen (INT 8), um Dokumente in unterschiedlichen Sprachen zu finden (INT 11), um Webseiten zu finden, die die Verwendung von Termini belegen (INT 11), oder um zusätzliche Quellen zu finden, auf Basis derer dann die Begriffsdefinitionen erstellt werden (INT 9). Auch der Termextraktionsschritt wurde von einigen erwähnt, sei es, dass sie Termextraktion durchführen (INT 9, INT 10, INT 12, INT 15), oder dass sie KI-basierte Termextraktionsmodule in Terminologieverwaltungssysteme integrieren werden (INT 1, INT 2). Als nächster Schritt im Terminologieworkflow käme nun die Erstel-

lung von terminologischen Einträgen. Keiner der Interviewteilnehmer*innen hat terminologische Einträge komplett von KI erstellen lassen, aber einzelne Aspekte, die zur Erstellung eines Eintrags beitragen, wurden mit KI getestet, darunter das Schreiben von Definitionen (INT 1, INT 2, INT 3, INT 6), die Übersetzung von Definitionen (INT 2, INT 6), die Suche nach Äquivalenten (INT 3, INT 5, INT 6, INT 7), die Suche nach der Etymologie (INT 5, INT 7), die Zuweisung eines Fachgebiets (INT 10, INT 11), das Erstellen von Begriffssystemen (INT 5). KI-Anwendungen wurden aber auch zur Begriffsabgrenzung (INT 7), zur Überprüfung von Übersetzungsäquivalenten (INT 11), für die Beschreibung von Begriffsbeziehungen (INT 5), für das Finden und Bestimmen von Synonymen (INT 11) und das Suchen von Abkürzungen (INT 11) erprobt. Für den Schritt der Revision und Validierung wurden KI-Anwendungen nicht so häufig getestet. Nur ein*e Interviewteilnehmer*in (INT 11) gab an, ChatGPT für die Revision einer Termliste getestet zu haben. Unter den Schritt Veröffentlichung und Verbreitung fallen z. B. auch das Konvertieren der Daten in andere Datenformate, was auch mit KI getestet wurde (INT 8) und im Rahmen der Wartung und Aktualisierung wurde die Bereinigung der Daten mit Hilfe von KI getestet (INT 10).

Obwohl fast alle Interviewteilnehmer*innen unterschiedliche Workflowschritte mit KI getestet haben, integrierten zum Zeitpunkt der Interviews nur wenige tatsächlich KI-Anwendungen bereits in ihre tägliche Terminologiearbeit (INT 10, INT 11, INT 12).

6.2 Rollen und Kompetenzen

6.2.1 Rollen

Im Moment liegt bei den meisten interviewten Expert*innen noch immer einer der Schwerpunkte auf der Erstellung von terminologischen Datensätzen und terminologischen Ressourcen. Die Validierung von von KI erstellten terminologischen Einträgen spielt noch keine große Rolle, da die automatische Erstellung eines terminologischen Eintrags, vor allem in Sprachen mit begrenzten Datenressourcen, noch nicht den Qualitätskriterien entspricht, wie das eine Interviewteilnehmer*in ausdrückt: „We don't consider the auton-

omous creation of a terminological record in [less-resourced language] to be satisfactory.“ (INT 8)

Es zeigt sich aber, dass sich die Schwerpunkte mit der Zeit verschieben können bzw. neue hinzukommen, wie das folgende Zitat aus dem Interviewkorpus unterstreicht:

„[A]ls ich begonnen habe, da ging es tatsächlich nur darum, Einträge zu machen, und jetzt kümmern wir uns halt um das Ganze, [...] wir kümmern uns um die Lizenzen für die elektronischen Wörterbücher, wir kümmern uns um das Tool, wir kümmern uns um die maschinelle Übersetzung und wahrscheinlich dann demnächst auch um ChatGPT und diese Geschichten.“ (INT 15)

Mit der Zeit wird sich zeigen, ob die Prognose, dass Terminologieexpert*innen immer mehr die Rolle von Terminologie-Kurator*innen annehmen werden (Heyn/Allard 2025) und somit nicht mehr selber Terminologieeinträge erstellen, sondern nur noch KI-generierte Einträge validieren werden, zutreffen wird. Im Interviewkorpus konnten dafür keine Hinweise gefunden werden.

6.2.2 Kompetenzen

In den Interviews wurden nicht direkt Kompetenzen abgefragt, aus den Interviewdaten lassen sich dennoch Kompetenzen ableiten. *Prompt Engineering* wurde in vielen Interviews erwähnt und auch wie zeitaufwendig es sein kann, einen Prompt zu erstellen, der das gewünschte Ergebnis liefert, wie dies dieses Zitat aus dem Interviewkorpus illustriert: „[I]n some cases it took several attempts to create a prompt that was clear enough for it to understand what exactly it was being asked to provide“ (INT 9). Daher ergibt es Sinn, dass Terminologieexpert*innen das Verfassen von Prompts erlernen. Auf der anderen Seite stellt sich die Frage, ob wirklich alle Personen, die im Terminologieworkflow beteiligt sind, selber Prompts schreiben können müssen, da z. B. Prompts auch zentral entwickelt werden können und dann von allen Mitarbeiter*innen in der Institution verwendet werden können, wie es dieses Zitat zeigt: „We managed to create a well-functioning prompt [for definition writing], which is currently available for all [staff] in the [internal AI] portal.“ (INT 10).

Eine wichtige Kompetenz ist es, die Grenzen der Sprachunterstützung zu erkennen, was dieses Zitat illustriert: „Since the [name of less-resourced language] is less developed in digital space and AI platforms, the proposed equivalents rarely provide acceptable options; however, they often serve as a source of inspiration for possible [name of less-resourced language] terms.“ (INT 5)

Eine weitere wichtige Kompetenz ist es, die Qualität, Aktualität und Validität des KI-Outputs zu überprüfen, wie dies durch dieses Zitat aus dem Interviewkorpus illustriert wird: „Ich habe ChatGPT getestet, um Vorschläge für Benennungsvarianten zu sammeln und Definitionen zu erstellen. Die Ergebnisse sind interessant und können als Inspiration dienen, müssen aber immer inhaltlich gegengeprüft werden.“ (INT 6)

Im Zusammenhang mit diesen beiden Kompetenzen spielen die bisherigen Kompetenzen aus dem Bereich der Terminologiarbeit und das Fachwissen aus dem Bereich, für den die Terminologie gerade erarbeitet wird, eine immer bedeutendere Rolle, da nur mit der Expertise aus diesen beiden Bereichen fehlerhafter KI-Output oder Halluzinationen erkannt werden können: z. B. wenn Termini erfunden werden oder wenn bei Vorschlägen zu neuen Benennungen die Regeln der Benennungsbildung nicht eingehalten werden.

Terminologieexpert*innen müssen auch in der Lage sein, die Anwendungsmöglichkeiten von KI im Workflow abzuwägen und nicht alle Arbeitsschritte zwangsweise mit KI lösen zu wollen. Terminologieexpert*innen sollten auch erklären und argumentieren können, dass andere Methoden vielleicht zielführender sind, wie dies ein*e Interviewteilnehmer*in beschreibt: „The possibility to use AI for [name of the termbase] cleaning was explored and the conclusions were that it was not the right way to go. It did not bring added value to the [existing] cleaning procedure and other means were more efficient.“ (INT 10)

In Bezug auf ethische und datenschutzrechtliche Überlegungen sollten Terminologieexpert*innen wissen, welche Daten sie sicher mit KI verarbeiten können und welche Daten nicht eingegeben werden sollten (Heinisch 2025). Das nachfolgende Zitat aus dem Interviewkorpus illustriert, dass Terminologieexpert*innen ethische und datenschutzrechtliche Überlegungen anstellen: „However, we seldom do this [using AI applications] because we are concerned about sharing information and how our data could be used.“ (INT 7) Weiters können durch die Aufklärung innerhalb der Institution und durch Verbreiten

von Informationen zu ethischen und datenschutzrechtlichen Überlegungen solche Kompetenzen aufgebaut werden: „[E]s ist allerdings schon Sensibilisierungsarbeit über einen Intranetartikel geleistet worden (keine Eingabe von personenbezogenen oder sensiblen Daten [...]).“ (INT 6)

7 Fazit und Ausblick

In dieser explorativen Studie wurde basierend auf qualitativen Interviews untersucht, welche Auswirkung KI auf Terminologearbeit im institutionellen Kontext haben kann, im Speziellen, welche Arbeitsschritte in ihr durch KI-Anwendungen unterstützt bzw. automatisiert werden können und wie sich dadurch die Rollen und die Kompetenzen von Terminologie*expert*innen verändern. Die explorative Studie hat gezeigt, dass in der Praxis sehr viel mit KI-Anwendungen experimentiert wird und dass sehr viele Schritte im Terminologieworkflow bereits durch KI-Anwendungen unterstützt werden können, dass aber zum Zeitpunkt der Interviews KI-Anwendungen noch nicht durchgängig in den Terminologieworkflow eingebunden waren und dass die Terminologieexpert*innen frei entscheiden konnten, KI-Anwendungen einzusetzen oder nicht. Weiters spielte die vollautomatisierte Erstellung von terminologischen Einträgen oder ganzen terminologischen Ressourcen mit Hilfe von KI noch keine Rolle.

Was die Rollenprofile betrifft, konzentrieren sich aktuell die meisten interviewten Terminologieexpert*innen vor allem auf die Erstellung terminologischer Einträge und Ressourcen. Die Überprüfung von KI-generierten terminologischen Einträgen ist bislang kaum relevant, da die automatisch erstellten Einträge, insbesondere bei ressourcenarmen Sprachen, noch nicht den gewünschten Qualitätsansprüchen genügen. Es wird sich zeigen, ob in Zukunft Terminologieexpert*innen zu Terminologiekurator*innen werden.

Bevor KI-Anwendungen allgegenwärtig waren, waren für Terminolog*innen Kompetenzen im Bereich der Terminologearbeit, Recherchekompetenz, Sachwissen und Sprachwissen wichtig. Mit der zunehmenden Digitalisierung wurde auch die Datenkompetenz wichtiger. Nun kommen die KI-Kompeten-

zen hinzu und die Datenkompetenz wird essentieller, da sie als Grundlage für die KI-Kompetenzen dient. Aber die etablierten Kompetenzen wie Kenntnisse der Terminologiewissenschaft und Methoden der Terminologiearbeit, die Sprachkenntnisse und das Sachwissen verlieren nicht an Bedeutung, sondern spielen eine immer bedeutendere Rolle, da nur mit der Expertise aus diesen Bereichen fehlerhafter KI-Output oder Halluzinationen erkannt werden können: z.B. wenn Termini erfunden werden, wenn bei Vorschlägen zu neuen Benennungen die Regeln der Benennungsbildung nicht eingehalten werden oder wenn Begriffsbeziehungen falsch oder Definitionen fehlerhaft sind.

Durch die rasche Entwicklung in diesem Bereich wäre es wünschenswert, die Terminologieexpert*innen im Abstand von ein oder zwei Jahren noch einmal zu interviewen, um zu sehen, ob die KI-Anwendungen nun bereits standardmäßig im Workflow integriert sind und ob sich die Rolle der Terminologiekurator*in manifestiert. Es wäre auch interessant, die Interviews mit Terminologieexpert*innen in Unternehmen durchzuführen und die Ergebnisse zu vergleichen.

Literatur

- CHIOCCHETTI, ELENA / LUŠICKY, VESNA / WISSIK, TANJA (2023): „Multilingual Legal Terminology Databases. Workflows and Roles.“ *Handbook of Terminology, vol 3: Legal Terminology*. Eds. Łucja Biel / Hendrik J. Kockaert. Amsterdam/Philadelphia: Benjamins. 458–484.
- DREWER, PETRA / MAYER, FELIX / PULITANO, DONATELLA (HRSG.) (2025): *Terminologie in der KI – KI in der Terminologie. Akten des Symposions*, Worms, 27.–29. März 2025. München: Deutscher Terminologie-Tag e. V.
- FLEISCHMANN, KLAUS / LANG, CHRISTIAN (2025): „Terminologie in der KI. Wie mit Terminologie der Output von LLMs und GenAI optimiert werden kann.“ Drewer/Mayer/Pulitano (2025): 83–95.
- GIACOMINI, LAURA (2024): „AI for Terminology Work: A Useful Device for Concept Modeling and the Creation of Domain Ontologies?“ Presentation held at *Concept Systems and Frames in Terminology (CSFT)*, Innsbruck, Austria, 24 September 2024.

- GIACOMINI, LAURA (2025): „Detecting Simple and Complex Collocations in the Technical Domain by Means of GenAI: How Far Can We Go?“ Presentation held at *EUROPHRAS 2025*. Aarhus, Denmark, 19–22 May 2025.
- GUEST, GREG / BUNCE, ARWEN / JOHNSON, LAURA (2006): „How Many Interviews are Enough? An Experiment with Data Saturation and Variability.“ *Field Methods* 18.1. 59–82.
- HEINISCH, BARBARA (2025a): „Terminologielehre – trotz oder mit Large Language Models?“ *Drewer/Mayer/Pulitano (2025)*: 121–133.
- HEINISCH, BARBARA (2025b): „Large Language Models for Terminology Work: A Question of the Right Prompt?“ *Journal for Language Technology and Computational Linguistics* 38.2: 13–30. DOI: 10.21248/jlcl.38.2025.280.
- HEYN, MATTHIAS / ALLARD, MARTA (2025): „Künstliche Intelligenz und Terminologie. Eine symbiotische Beziehung.“ *Drewer/Mayer/Pulitano (2025)*: 61–81.
- HIBBERT, MELANIE / ALTMANN, ELENA / SHIPPEN, TRISTAN / WRIGHT, MELISSA (2024): „A Framework for AI Literacy.“ <https://academiccommons.columbia.edu/doi/10.7916/cvjjz-ve81> (13.11.2025).
- KHEMAKHEM, MOHAMED / VALENTINI, CRISTINA / RALLI, NATASCIA / BARROS, SÉRGIO / LÖCKINGER, GEORG / VEZZANI, FEDERICA / SALGADO, ANA / ZHANG, ZHENLING / MAHR, SABINE / CARVALHO, SARA / FLEISCHMANN, KLAUS / COSTA, RUTE (2025): „Terminology Management Meets AI: The ISO/TC 37/SC 3/WG 6 Initiative.“ *Proceedings of TermTrends 2025 Workshop at LDK*, University of Naples „L’Orientale“, Sep 2025, Naples, Italy. <https://hal.science/hal-05162487v2> (12.12.2025).
- KUCKARTZ, UDO (2014): *Qualitative Text Analysis. A Guide to Methods, Practice and Using Software*. London: SAGE Publications.
- VAN DER LEK, IULIANNA / WOLDRICH, ANNA / WISSIK, TANJA (2025): „AI in Higher Education: A CLARIN Community Survey.“ *Digital Humanities in the Nordic and Baltic Countries Publications* 7.2. DOI: 10.5617/dhnbpub.12303.
- LERNER, PAUL / YVON, FRANÇOIS (2025): „Towards the Machine Translation of Scientific Neologisms.“ *Proceedings of the 31st International Conference on Computational Linguistics*. Abu Dhabi, UAE, 947–963. <https://aclanthology.org/2025.coling-main.63.pdf> (29.11.2025).
- LEWIS, PATRICK / PEREZ, ETHAN / PIKTUS, ALEKSANDRA / PETRONI, FABIO / KARPUKHIN, VLADIMIR / GOYAL, NAMAN / KÜTTLER, HEINRICH / LEWIS, MIKE /

- YIH, WEN-TAU / ROCKTÄSCHEL, TIM / RIEDEL, SEBASTIAN / KIELA, DOUWE (2020): „Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks.“ *Proceedings of Advances in Neural Information Processing Systems* 33 (NeurIPS 2020). <https://proceedings.neurips.cc/paper/2020/file/6b493230205f780e1bc26945df7481e5-Paper.pdf> (13.11.2025).
- LONG, DO XUAN / DINH, DUY / NGUYEN, NGOC-HAI / KAWAGUCHI, KENJI / CHEN, NANCY F. / JOTY, SHAFIQ / KAN, MIN-YEN (2025): „What Makes a Good Natural Language Prompt?“ *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (ACL) (Volume 1: Long Papers)*, Vienna, Austria. Association for Computational Linguistics. 5835–5873. <https://aclanthology.org/2025.acl-long.292/> (13.11.2025).
- LONG, DURI / MAGERKO, BRIAN (2020): „What is AI Literacy? Competencies and Design Considerations.“ *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems (CHI '20)*. Association for Computing Machinery, New York, NY, USA. 1–16. DOI: 10.1145/3313831.3376727.
- ŁUKASIK, MARKE (2023): „Terminology Work in the Era of AI-generated Content. A Case of ChatGPT.“ Presentation held at *5th International Conference on Terminology: Scientific, Administrative and Educational Dimensions of Terminology*. 19–20 Oct. 2023. https://www.researchgate.net/publication/378297882_Terminology_work_in_the_era_of_AI-generated_content_A_case_of_ChatGPT_an_end-user_perspective (13.11.2025).
- MASSION, FRANÇOIS (2024): „Terminology in the Age of AI: The Transformation of Terminology Theory and Practice.“ *Journal of Translation Studies* 4.1: 67–94. DOI: 10.3726/JTS012024.04.
- MEUSER, MICHAEL / NAGEL, ULRIKE (1991): „Experteninterviews – vielfach erprobt, wenig bedacht.“ *Qualitative-empirische Sozialforschung. Konzepte, Methoden, Analysen*. Eds. Detlef Gerz / Klaus Karaimer. Opladen: Westdeutscher Verlag. 441–471.
- NAGENDRA, NAGARAJ / CHANDRA, JAYARAMAN (2022): „A Systematic Review on Features Extraction Techniques for Aspect Based Text Classification Using Artificial Intelligence.“ *ECS Transactions* 107.1: 2503.
- RADT (2020): *Berufsprofil 2.0. Terminologin – Terminologe*. https://publikationen.radt.org/RaDT_Berufsprofil_2-0_2020_deutsch.pdf (13.11.2025).

- RATISLAVOVÁ, KATEŘINA / RATISLAV, JAKUB (2014): „Asynchronous Email Interview as a Qualitative Research Method in the Humanities.“ *Human Affairs* 24: 452–460. DOI: 10.2478/s13374-014-0240-y.
- REINEKE, DETLEF (2023): „Terminologiewerk mit ChatGPT & Co.“ *Fachzeitschrift für Terminologie* 19.1: 25–28. <https://edition.dttv.org/ausgaben/edition-2023-1-e-version.pdf> (13.11.2025).
- SAN MARTÍN, ANTONIO (2024): „What Generative Artificial Intelligence Means for Terminological Definitions.“ *Proceedings of the 3rd International Conference on Multilingual Digital Terminology Today (MDTT 2024)*, Granada, Spain, June 27–28. <https://ceur-ws.org/Vol-3703/paper1.pdf> (13.11.2025).
- SÁNCHEZ-GIJÓN, PILAR / PALENZUELA-BADIOLA, LEIRE (2023): „Analysis and Evaluation of ChatGPT-Induced HCI Shifts in the Digitalised Translation Process.“ *Proceedings of the International Conference on Human-informed Translation and Interpreting Technology 2023*. 227–267. DOI: 10.26615/issn.2683-0078.2023_021.
- SCHÜLLER, KATHARINA / BUSCH, PAULINA / HINDINGER, CARINA (2019): *Future Skills: Ein Framework für Data Literacy – Kompetenzrahmen und Forschungsbericht. Arbeitspapier Nr. 47*. Berlin: Hochschulforum Digitalisierung. DOI: 10.5281/zenodo.3349865.
- SINGHAL, KARAN / TU, TAO / GOTTSWEIS, JURAJ / SAYRES, RORY / WULCZYN, ELLERY / AMIN, MOHAMED / HOU, LE / CLARK, KEVIN / PFOHL, STEPHEN R. / COLE-LEWIS, HEATHER / NEAL, DARLENE / RASHID, QAZI MAMUNUR / SCHAEKERMANN, MIKE / WANG, AMY / DASH, DEV / CHEN, JONATHAN H. / SHAH, NIGAM H. / LACHGAR, SAMI / MANSFIELD, PHILIPP ANDREW / PRAKASH, SUSHANT / GREEN, BRADLY / DOMINOWSKA, EWA / AGÜERA Y ARCAS, BLAISE / TOMAŠEV, NEED / LIU YUN / WONG, RENEE / SEMTURS, CHRISTOPHER / MAHDAVI, SARA S. / BARRAL, JOELLE K. / WEBSTER, DALE R. / CORRADO, GREG S. / MATIAS, YOSHI / AZIZI, SHEKOOFEH / KARTHIKESALINGAM, ALAN / NATARAJAN, VIVEK (2025): „Toward Expert-level Medical Question Answering with Large Language Models.“ *Nature Medicine* 31: 943–950. DOI: 10.1038/s41591-024-03423-7.
- STRAUB, STEFAN (2024): *CAIL – Critical AI Literacy (Endbericht)*. Vienna. <https://epub.oeaw.ac.at/ita/ita-projektberichte/ITA-2024-07.pdf> (13.11.2025).
- TERMINORGS (2016): *Recommended Skills for Terminology Professionals*. https://www.terminorgs.net/downloads/Terminology_skills.pdf (13.11.2025).

- VAN DER LEK, IULIANNA / WOLDRICH, ANNA / WISSIK, TANJA (2025): „AI in Higher Education: A CLARIN Community Survey.“ *Digital Humanities in the Nordic and Baltic Countries Publications* 7.2. DOI: 10.5617/dhnbpub.12303.
- VARGA, CHRISTINA (2024): „El uso de ChatGPT en el trabajo terminológico puntual y sistemático en rumano y en castellano.“ *Quaestiones Romanicae* XI: 377–396. DOI: 10.35923/QR.11.02.32.
- WEBERSINKE, NICOLAS / KRAUS, MATHIAS / BINGLER, JULIA ANNA / LEIPPOLD, MARKUS (2021): „ClimateBert: A Pretrained Language Model for Climate-Related Text.“ ArXiv abs/2110.12010 (2021) (13.11.2025).
- WISSIK, TANJA (2025a): „Impact of Automatic Term Extraction on Terminology Work. A Qualitative Interview Study in Institutional Settings.“ *Terminology* 31.1: 110–136.
- WISSIK, TANJA (2025b): „An Infrastructural Approach to Terminology Work.“ *Selected Papers of the CLARIN Annual Conference 2024*. Eds. Vincent Vandeghinste / Thassia Kontino. Linköping Electronic Conference Proceedings. 1–12.
- WISSIK, TANJA (2025c): „The Use of Translation Technologies and Generative Artificial Intelligence by Terminology Professionals in Institutional Settings: A Case Study.“ *tradumatica* 2025: 165–192.

Terminologie und Retrieval-Augmented Generation (TermRAG)

Abstract

In this article, we introduce the concepts of Retrieval-Augmented Generation (RAG) and TermRAG. RAG is a technique that combines the capabilities of generative language models with information-retrieval methods for more precise and reliable output. TermRAG is a specialized variant of RAG that leverages terminological resources during the retrieval process. We further discuss application scenarios of TermRAG in specialized communication and outline the current state of technical implementation options.

Keywords generative language models, RAG, specialized communication, terminological resources

1 Konzept RAG und TermRAG

Der rasante Einzug der KI in den Alltag und in die Arbeitswelt ist in letzter Zeit vor allem generativen Sprachmodellen¹ zu verdanken, die durch ihre einfache

1 In diesem Beitrag verwenden wir die Bezeichnung „generative Sprachmodelle“ für Transformer-basierte (Vaswani et al. 2017), autoregressive *Decoder-only*-Sprachmodelle (Jurafsky/Martin 2025: Kap. 7, 8), die In-Context-Learning-Verhalten zeigen. Üblich ist ebenfalls die Bezeichnung LLM (Large Language Model); da es aber innerhalb der generativen Sprachmodelle, die In-Context-Learning-Verhalten zeigen, erhebliche Größenunterschiede gibt (z. B. einstellige Milliardenparameteranzahl bis hin zu Billionenparameteranzahl), entscheiden wir uns für die neutrale Bezeichnung „generative Sprachmodelle“.

Bedienbarkeit, Anpassbarkeit und beeindruckende Fähigkeit zur Sprachverarbeitung und -generierung neue Möglichkeiten eröffnen.

Auch in der ein- und mehrsprachigen Fachkommunikation ist der Einsatz allgemeiner generativer Sprachmodelle möglich. Allerdings verfügen solche Modelle ohne weitere Anpassung in der Regel nicht über das notwendige spezifische (Fach-)Wissen, obgleich sie mit immensen Datenmengen trainiert werden (Krüger 2025: 13; Zhao et al. 2023: 17f.). Sollte ein Unternehmen ein generatives Sprachmodell zur Erstellung von Fachtexten, beispielsweise Informationsprodukten (DIN 82079: 2021) oder ihren Übersetzungen, benutzen, fehlen dem Modell typischerweise aktuelle Unternehmensdaten wie die aktuelle Produktpalette oder firmeninterne terminologische Vorgaben, sofern relevante Informationen nicht öffentlich waren oder erst nach Abschluss des Pre-Trainings veröffentlicht wurden. Aber selbst bei öffentlich verfügbaren Informationen ist aufgrund der probabilistischen Funktionsweise generativer Sprachmodelle sowie der Vielzahl und Heterogenität der Trainingsdaten nicht gewährleistet, dass bei der Textgenerierung konsistent, vollständig und exakt die intendierten fachlichen Inhalte reproduziert werden.

Allerdings zeigen generative Sprachmodelle ab einer bestimmten Größe ein sog. In-Context-Learning-Verhalten (ICL) (Krüger 2025: 14; Zhao et al. 2023: 3): Durch das Präsentieren von Beispielen und/oder die Bereitstellung detaillierter Informationen im Prompt können sie in speziellen Kontexten spezifische Aufgaben übernehmen, für die sie nicht explizit trainiert wurden. ICL kann dafür genutzt werden, dem Sprachmodell externe, also nicht in den Trainingsdaten vorhandene Informationen wie zum terminologischen Gebrauch oder zu Produkten zu präsentieren, die es dann (etwa für eine Terminologieprüfung) bei der Antwortgenerierung nutzen kann.

Eine einfache Möglichkeit, kleinere Datensammlungen für Sprachmodelle bereitzustellen, ist ihre vollständige Integration in den Prompt. Für größere Sammlungen gibt es häufig die Option zum Datei-Upload. Beide Möglichkeiten sind leicht umzusetzen, haben aber mehrere Nachteile. Zum einen muss bei der Nutzung unternehmensexterner Dienste abgewogen werden, ob aus Gründen der Datensouveränität die komplette Datensammlung per Prompt an ein Sprachmodell gesendet bzw. auf einen Server hochgeladen werden soll.

Zum Zweiten besteht ein weiterer Nachteil der vollständigen Datenintegration darin, dass damit das sog. Kontextfenster des Sprachmodells erreicht werden kann. Als Kontextfenster gilt die Menge an Tokens, die ein Modell in einer Aufgabe gleichzeitig berücksichtigen kann; dazu zählen nicht nur die Informationen aus dem Prompt, sondern ggf. auch vorherige Konversationen. Muss ein Modell viele irrelevante Informationen verarbeiten (vgl. Li et al. 2024: 881), stellt sich zusätzlich die Frage nach einem ressourcenschonenden Einsatz generativer KI.

Schließlich muss bei einer vollständigen Prompt-Integration von eigentlich strukturierten Daten als unstrukturierter Text sichergestellt werden, dass das Sprachmodell die zusammenhängenden Informationseinheiten richtig erkennt: Bei begriffsorientierten terminologischen Daten soll es die Begriffsgrenzen und somit beispielsweise die Zugehörigkeit einer Benennung samt Verwendungsstatus zu einem Begriff korrekt erkennen. Dieser Problematik kann man versuchen durch Prompt-Engineering (vgl. Schulhoff et al. 2024) zu begegnen.

Eine Technik, bei der einem Sprachmodell hingegen zielgerichtet relevante Zusatzinformationen aus einer externen Wissensbasis zusätzlich zum Prompt zur Verfügung gestellt werden, ist die *Retrieval-Augmented-Generation-Technik* (RAG-Technik) (Lewis et al. 2021; Gao et al. 2024).

Die allgemeine RAG-Arbeitsweise sieht wie folgt aus. Der User gibt einen Prompt ein, der aber nicht direkt an das Sprachmodell geschickt wird. Stattdessen wird zuerst eine Wissensbasis nach Informationen durchsucht, die für den Prompt relevant sind. Dabei werden Informationen im Prompt mit den Informationen in der Wissensbasis verglichen und die Relevanzbewertung erfolgt mit geeigneten Methoden des Information Retrievals (bspw. Stock/Stock 2015). Damit diese Relevanzbewertung effizient erfolgen kann, wird die Wissensbasis im Vorfeld so aufbereitet (Indexierung), dass gezielt und schnell auf passende Inhalte zugegriffen werden kann. Sind die relevanten Informationen identifiziert, wird der Prompt um diese Informationen ergänzt und man spricht dann von einem augmentierten Prompt. Anschließend wird der augmentierte Prompt an das Sprachmodell gesendet. Das Sprachmodell verarbeitet den augmentierten Prompt und generiert Output.

Im Namen der RAG-Technik verbergen sich die drei zentralen Verarbeitungsschritte:

Retrieval – umfasst den Schritt der Identifizierung relevanter Abschnitte in der Wissensbasis.

Augmentierung – umfasst den Schritt der Ergänzung des Prompts um die selektierten Abschnitte aus der Wissensbasis.

Generation – umfasst den Schritt der Verarbeitung des augmentierten Prompts durch das Sprachmodell und die Generierung des entsprechenden Outputs.

RAG bezeichnet also die Technik, bei der ein generatives Sprachmodell um eine oder mehrere anwendungsspezifische Wissensbasen ergänzt wird. Dabei wird das Sprachmodell weder mit den anwendungsspezifischen Daten neu trainiert noch verarbeitet es die komplette Wissensbasis. Stattdessen benutzt das Sprachmodell zur Generierung einer Antwort einen um selektierte Informationen augmentierten Prompt.

Eine Möglichkeit, die RAG-Technik umzusetzen, ist es, über bestehende Suchmaschinen auf das Internet als (bereits indexierte) externe Wissensbasis zuzugreifen und in Echtzeit nach relevanten Informationen für die jeweilige Aufgabe zu suchen. Neben urheber- und datenschutzrechtlichen Aspekten (Zhong et al. 2025) und dem Umstand, dass auf Informationen zugegriffen wird, die außerhalb des eigenen Einflusses liegen, ist ein solches Vorgehen auch nicht geeignet, nicht-öffentliche Zusatzinformationen zu liefern. Daher liegt nahe, eigene Datensammlungen als Wissensbasis zu nutzen. Denn viele Unternehmen verfügen über kuratierte, spezifische Datensammlungen, wie Terminologiedatenbanken, die aufwändig über die Jahre aufbereitet wurden, aber nicht öffentlich zugänglich sind. Auch diese Datensammlungen können Sprachmodellen für unternehmensspezifische fachkommunikative Aufgaben bereitgestellt werden.

TermRAG ist eine spezielle Art der RAG-Technik, bei der terminologische Daten als Wissensbasen in RAG-Workflows benutzt werden, wobei dafür in

der Community auch weitere Bezeichnungen wie TAG (Fleischmann/Lang 2025; Lackner et al. 2025a, b; Nunzio 2025) oder M-RAG (Hummer 2023) existieren.

2 Einfache TermRAG-Pipeline und Anpassungsmöglichkeit

Es gibt verschiedene Möglichkeiten, eine TermRAG-Pipeline zu konzeptionieren und basierend darauf konkret zu implementieren. Im Folgenden wird eine einfache TermRAG-Pipeline konzeptionell skizziert, die sich an allgemeinen Naive RAG-Pipelines orientiert (bspw. Gao et al. 2024: 3), aber nicht für die Arbeit mit terminologischen Daten optimiert ist. Anschließend werden die aktuellen Anpassungsmöglichkeiten dieser einfachen TermRAG-Pipeline entlang der Verarbeitungsschritte diskutiert.

Für die Indexierung können in einer einfachen TermRAG-Pipeline zunächst Inhalte aus einer Terminologiedatenbank exportiert werden, beispielsweise als Textdatei, in der alle terminologischen Einträge untereinander als fortlaufender Text gespeichert sind. Danach wird diese Textdatei in kleinere Abschnitte (sog. Chunks) segmentiert. Die Segmentierung kann nach einer festen Anzahl von Tokens, z. B. 100 Tokens, erfolgen, wobei eine Überlappung zwischen den Chunks, z. B. 20 Tokens, vorgesehen wird. In dieser einfachen Pipeline kommt semantisches Retrieval zum Einsatz, das auf Vektorenähnlichkeit beruht. Dafür müssen die Chunks vektorisiert, also in eine Vektorrepräsentation überführt werden. Die so erzeugten Vektoren werden in einer Vektordatenbank indexiert (Wissensbasis), sodass später über Ähnlichkeitssuche relevante Chunks abgerufen werden können. Die Vektorisierung erfolgt mithilfe eines Embedding-Modells.

Im Retrieval-Schritt wird der eingegebene (User-)Prompt ebenfalls in eine Vektorrepräsentation überführt. Diese wird mit den Vektorrepräsentationen der Chunks in der Wissensbasis verglichen. Dabei kann als Ähnlichkeitsmaß die Cosinus-Ähnlichkeit verwendet werden. Die ähnlichsten Chunk-Vektoren

zum gegebenen Prompt-Vektor werden für den nächsten Schritt der Augmentierung bereitgehalten.

Im Augmentierungsschritt wird der Inhalt der abgerufenen Chunks als zusätzlicher Text in den ursprünglichen Prompt integriert. Der so augmentierte Prompt wird im Generierungsschritt vom Sprachmodell zu einem Output verarbeitet.

Bei dieser einfachen TermRAG-Pipeline fallen dem terminologisch geschulten Auge sofort Schwachstellen auf: Beispielsweise wird mit terminologischen Daten in Textform gearbeitet, obwohl ein strukturiertes Format wie XML naheliegender ist. Deshalb kann diese einfache TermRAG-Pipeline für die Arbeit mit terminologischen Daten angepasst und optimiert werden. Im Folgenden wird der aktuelle Stand der Anpassungsmöglichkeiten entlang aller Verarbeitungsschritte näher erläutert.

2.1 Indexierung

In der skizzierten TermRAG-Pipeline wird mit einer exportierten Textdatei gearbeitet, was aber, wie erwähnt, nicht optimal ist. Für Terminologiedatenbanken bieten sich insbesondere strukturierte Formate, beispielsweise XML (bspw. TBX [ISO 30042: 2019] oder Tool-spezifische XML-Dialekte), aber auch JSON und RDF(S)-basierte Graphformate (Gao et al. 2024: 7) an. In aktuellen Arbeiten zu TermRAG werden verschiedene Formate genutzt (Lackner et al. 2025a, b; Lang et al. eingereicht) und ihre Eignung wird systematisch untersucht (Lackner et al. 2025a, b).

Der nächste Aspekt betrifft die Chunking-Strategie, also die Aufbereitung der Daten für die spätere Indexierung. In der skizzierten TermRAG-Pipeline wurde das Chunking nach einer festen Tokenanzahl gewählt. Dieses Vorgehen kann für eine begriffsorientierte Terminologiedatenbank dazu führen, dass Informationen, die zu einem Begriff gehören, auf zwei Chunks aufgeteilt werden oder dass Informationen zu mehreren Begriffen in einem Chunk enthalten sind. Daher bietet sich für begriffsorientierte Terminologiedatenbanken das Chunken nach Begriffen an (Lackner et al. 2025a; Lang et al. eingereicht). Bei dieser Chunking-Strategie repräsentiert ein Chunk einen Begriff und im

Retrieval-Schritt kann folglich nach relevanten Begriffen (einschließlich weiterer Informationen zu diesem Begriff auf Sprach- und Benennungsebene) gesucht werden.

Ist die Chunking-Strategie festgelegt, muss geprüft werden, ob für den Retrieval-Prozess die Daten aus der ursprünglichen Repräsentation in eine andere Repräsentation überführt werden sollten, um die angedachte Retrieval-Methode anwenden zu können. In der skizzierten TermRAG-Pipeline werden die Chunks mit einem Embedding-Modell vektorisiert. Als Optimierungsschritt können verschiedene Embedding-Modelle getestet werden. Das Retrieval muss nicht zwingend auf Vektorrepräsentationen basieren (siehe unten); die Chunks können daher ggf. auch im ursprünglichen Format (XML, JSON, RDF(S), Text) belassen werden.

Neben der Arbeit mit exportierten Dateien besteht schließlich auch die Möglichkeit von direkten Datenbankabfragen aus dem Terminologieverwaltungssystem, beispielsweise über einen API-Zugriff (Fleischmann/Lang 2025) oder das Integrationsframework MCP (Anthropic 2025; vgl. Ferrari/Fleischmann 2025). Ein direkter, dynamischer Datenbankzugriff hat den Vorteil, dass terminologische Daten zur Erstellung der Wissensbasis stets aktuell sind; das Terminologieverwaltungssystem muss diese Zugriffsmethoden allerdings unterstützen.

2.2 Retrieval

In der skizzierten TermRAG-Pipeline beruht das Retrieval auf der Berechnung von Ähnlichkeiten zwischen dem Vektor des aktuellen Prompts und den Vektoren der Chunks mithilfe der Cosinus-Ähnlichkeit. Als Optimierungsschritt können andere Ähnlichkeitsmaße für Vektoren getestet werden (Hazledine 2024; You 2025).

Wie erwähnt, ist es nicht zwingend notwendig, vektorenbasierte Retrieval-Methoden zu verwenden. Für terminologische Daten in einfacheren Anwendungsszenarien, z. B. zur Terminologiekontrolle (s. unten), eignen sich beispielsweise auch zeichenkettenbasierte Retrieval-Methoden, bspw. die Levenshtein-Distanz (Stock/Stock 2015: 233) oder Fuzzy-Matching (Lackner et

al. 2025a: o.S.), wobei hier die bekannten Schritte des linguistischen Preprocessing wie Stemming oder Lemmatisierung berücksichtigt werden können.

Die Arbeit mit terminologischen Daten in einem graphbasierten Format bringt weitere Möglichkeiten des semantischen Retrievals mit sich, bei denen semantische Relationen genutzt werden können, um bspw. zusammenhängende Chunks zu identifizieren. Die graphbasierten Retrieval-Methoden können zur Ergänzung der beiden oben beschriebenen Retrieval-Strategien genutzt werden: Wurden die relevanten Chunks vektor- oder zeichenkettenbasiert identifiziert, können mithilfe semantischer Abfragen weitere relevante Chunks regelbasiert identifiziert werden (vgl. Baldassare 2025). Verfügt die Terminologiedatenbank über ein Begriffssystem, so können zu einem als relevant identifizierten Begriff regelbasiert beispielsweise alle seine Unterbegriffe oder Oberbegriffe zur Prompt-Augmentierung ebenfalls abgerufen werden.

2.3 Augmentierung

Bei der Augmentierung werden die Inhalte der relevanten Chunks als Text an den Prompt angehängt. Zur Optimierung kann die Anzahl der Chunks, die zur Augmentierung genutzt werden, angepasst werden. Darüber hinaus kann erwogen werden, ob der Prompt neben dem Chunk-Inhalt um weitere Angaben erweitert werden soll. Für terminologische Datenkategorien kann beispielsweise nur der Inhalt der Datenkategorie (z. B. der eigentliche Definitionstext) in den Prompt integriert werden.

Alternativ können dem Inhalt der Name der Datenkategorie („Definition:“) oder eine natürlichsprachliche Phrase vorangestellt werden (beispielsweise: „Der Begriff wird wie folgt definiert:“), um die Bedeutung der Zusatzinformation noch expliziter zu machen. Hierbei kann überlegt werden, ob diese Zusatzinformationen bereits bei der Erstellung der Wissensbasis enthalten sein sollen, damit sie im Retrieval genutzt werden können (vgl. Lang et al. eingereicht).

2.4 Generierung

Die Anpassungsmöglichkeiten im Generierungsschritt umfassen insbesondere die Wahl des Sprachmodells und das Prompt-Engineering, auf die hier allgemein eingegangen wird.

Bei der Wahl des Sprachmodells ist neben genereller Verfügbarkeit insbesondere seine Größe ein Faktor. Eine Überlegung ist dabei, dass in Kombination mit einer (Term)RAG-Pipeline bereits kleinere Sprachmodelle, beispielsweise mit einer einstelligen Milliardenanzahl Parameter, eingesetzt werden könnten, was Vorteile im Sinne der Hardwareanforderung, Nutzung auf lokalen Infrastrukturen und Ressourcenschonung mit sich bringen könnte. Die ersten Studien zeigen für TermRAG allerdings, dass für bestimmte Anwendungsszenarien, insbesondere für das Wissensmanagement und Question Answering, größere Sprachmodelle deutlich bessere Ergebnisse liefern (Lang et al. eingereicht). Weitere Untersuchungen zur Nutzung von kleineren Sprachmodellen im TermRAG-Kontext sind notwendig.

Der zweite Generierungsaspekt betrifft das Prompt-Engineering. Die wichtigste Aufgabe besteht darin, das Sprachmodell durch passendes Prompting so zu steuern, dass es bei der Generierung (fast) ausschließlich auf die aus der Wissensbasis eingeholten Informationen zugreift und dadurch Halluzinationen vermeidet. Dazu existieren bereits RAG-spezifische Prompt-Vorlagen (bspw. Langchain o.J.), die für eine TermRAG-Anwendung angepasst und weiterentwickelt werden können.

Die geschilderten Anpassungsmöglichkeiten einer terminologiespezifischen RAG-Pipeline hängen mit der konzeptionellen und technischen Beschaffenheit vorliegender terminologischer Daten sowie mit technischen Möglichkeiten zur RAG-Implementierung zusammen. Den entscheidenden Faktor stellt aber das Anwendungsszenario und seine Komplexität dar. Im Folgenden werden deshalb Anwendungsszenarien für TermRAG skizziert.

3 Anwendungsszenarien

Die Einsatzmöglichkeiten der RAG-Technik umfassen die traditionellen Anwendungsszenarien für präskriptive Terminologiedatenbanken, gehen aber auch darüber hinaus. Insbesondere können dem Sprachmodell mithilfe der RAG-Technik Vorgaben zu bevorzugten und verbotenen Benennungen für die ein- und mehrsprachigen Textproduktion bereitgestellt werden. Dieser Einsatzbereich wird traditionell in zwei Anwendungsszenarien unterteilt: Terminologiekontrolle, meistens aus der einsprachigen Perspektive, und die Erstellung von terminologisch korrekten (Fach-)Übersetzungen.

Eine andere Möglichkeit umfasst die Nutzung von terminologischen Daten als Quellen von anwendungsspezifischem, z. B. unternehmensspezifischem, Wissen über die Sachverhalte und deren Zusammenhänge im Sinne des Wissensmanagements. Im Folgenden wird die RAG-Nutzung in den drei Anwendungsszenarien skizziert.

3.1 Terminologiekontrolle

Beim ersten Anwendungsszenario handelt es sich um die Kontrolle der Terminologieverwendung in Fachtexten, bei dem kein Sprachwechsel im Sinne einer interlingualen Übersetzung stattfindet. Diese Kontrolle kann entweder zum Zeitpunkt der Texterstellung oder nachträglich erfolgen (Drewer/Ziegler 2014: 219–286). In der praktischen Terminologiarbeit war die nachträgliche Kontrolle bislang mithilfe von regelbasierten Tools wie beispielsweise *Controlled Language Checker* oder Terminologiechecker etabliert (Drewer/Ziegler 2014: 227–268).

Diese Tool-gestützte Terminologiekontrolle besteht gewöhnlicherweise aus zwei Schritten: der Textüberprüfung und einer anschließenden Textkorrektur. Der Einsatz von TermRAG-basierten Anwendungen zur Terminologiekontrolle kann die bekannten Schwächen regelbasierter Systeme insbesondere beim Korrekturschritt beheben (Gillhausen et al. 2024: 35): Bei der automatischen morphosyntaktischen Anpassung, wenn beispielsweise eine Benennungsänderung eine Artikelanpassung im Satz nach sich zieht, sowie bei der Disambigu-

ierung mehrdeutiger Benennungen. Darüber hinaus gibt es hybride Ansätze, die regelbasierte und TermRAG-gestützte Terminologiekontrolle kombinieren (Gillhaußen et al. 2024).

3.2 Terminologisch korrekte Textübersetzung

Das zweite Anwendungsszenario – Erstellung von terminologisch korrekten Übersetzungen – unterscheidet sich von der oben beschriebenen Terminologiekontrolle insofern, als hier ein Sprachwechsel stattfindet: Ausgehend von der Terminologieverwendung im Ausgangstext erfolgt eine terminologisch korrekte Übersetzung in die Zielsprache.

Für maschinelle Übersetzung werden aktuell sowohl speziell trainierte neuronale maschinelle Übersetzungssysteme als auch allgemeine generative Sprachmodelle eingesetzt und die Einbindung von terminologischen Vorgaben kann zu drei Zeitpunkten stattfinden: Beim Training der Modelle, während des Übersetzungsprozesses oder beim Postediting (Nesbigall 2025). Eine TermRAG-Pipeline wird insbesondere beim Postediting eingesetzt, um einen allgemein übersetzten Text an die terminologischen Vorgaben anzupassen (Hamm 2025). Gegenüber anderen Anbindungsmöglichkeiten von Terminologievorgaben wie Glossaren (Reineke 2025) profitiert ein Sprachmodell von einer TermRAG-Pipeline dadurch, dass diese nicht nur reine benennungsbezogene Informationen, sondern auch weitere Kontextinformationen in Definitionen oder weiteren Feldern bereitstellt, die auch in diesem Anwendungsszenario bei einer Disambiguierung der Termini helfen können (Hamm 2025).

3.3 Wissensmanagement

Wie schon in den beiden anderen Anwendungsszenarien deutlich wurde, ist eine Terminologiedatenbank nicht nur eine Quelle von Vorzugsbenennungen, sondern sie liefert durch ihre begriffsorientierte Struktur insbesondere auch eine unternehmensspezifische Sicht auf relevante Sachverhalte und ihre Zusammenhänge. Terminologiedatenbanken können deshalb im Sinne des

Wissensmanagements (Lang/Suchowolec 2020) in Auskunftssystemen wie KI-Chatbots eingesetzt werden, um beispielsweise Definitionsfragen zu beantworten. Als Wissensmanagement kann auch die allgemeinere Nutzung von TermRAG-basierten Anwendungen zur Generierung von Texten aufgefasst werden, die nicht nur auf Benennungs-, sondern auch auf Begriffsebene unternehmensspezifisch sein sollten (z. B. Mails, die die Produktpalette beschreiben [Ferrari/Fleischmann 2025]). Dieses Anwendungsszenario ist deshalb möglich, weil ein Sprachmodell Informationen zu fachlichen Zusammenhängen insbesondere im Feld „Definition“, aber auch auf der Benennungsebene im Feld „Kontext“ einer (präskriptiven) Terminologiedatenbank findet (vgl. Lang et al. eingereicht). Weiter gefasste Terminologiedatenbanken wie das Eco-Lexicon (Reimerink et al. 2010) enthalten häufig mehr enzyklopädische Informationen sowohl in Definitionen als auch in weiterführenden Kommentar- und Erklärungsfeldern sowie Verweise auf relevante Quellen im Internet, die noch mehr anwendungsspezifische Informationen für das Sprachmodell liefern können. Die ersten Studien zu TermRAG-basierten Chatbots zeigen, dass eine präskriptive Terminologiedatenbank bei diversen Fragen einen Ansatzpunkt liefern kann (Lang et al. eingereicht). Dabei ist allerdings zu beachten, dass eine Terminologiedatenbank in praktischen Wissensmanagementanwendungen nicht alleinige Wissensquelle sein kann, da sie durch ihre sprachliche Ausrichtung bei der Vielfalt an Fragestellungen an ihre Grenzen stößt.

4 Ausblick

TermRAG ist eine Technik, die neue Möglichkeiten einer kontrollierten Bereitstellung terminologischer Ressourcen an generative Sprachmodelle bietet. Dabei soll TermRAG von anderen Möglichkeiten der Bereitstellung wie einer direkten und vollständigen Integration der Terminologievorgaben im Prompt sowie der Promptmodifikation durch direkte Ersetzung nicht erlaubter Terminologie im Prompttext unterschieden werden. Deshalb sind für TermRAG im Unterschied zu anderen Ansätzen insbesondere der Aspekt eines modell-externen Retrievals und der Augmentierung ausschlaggebend.

TermRAG wird aktuell in der Forschung und Praxis intensiv erprobt und untersucht, sodass bereits kurzfristig mit neuen Entwicklungen zu rechnen ist. Neben einer systematischen Auswertung von Optimierungsmöglichkeiten für TermRAG-Pipelines stellt eine systematische Wirksamkeitsuntersuchung des TermRAG-Ansatzes, auch im Unterschied zu anderen Ansätzen wie regelbasierten Anwendungen oder *Long Context Language Models* (Li et al. 2024; Xu et al. 2024), ein wichtiges Desideratum dar. Dafür müssen die Evaluationsmethoden für TermRAG-generierte Texte noch stärker in den Fokus rücken, um die Erkenntnisse der bereits vorhandenen Evaluationsstudien und die dort eingesetzten Metriken (Blue und COMET, vgl. Lackner et al. 2025a, b, LLM-as-Judge und menschliche Evaluation [Lang et al. eingereicht]) systematisch auszuwerten und zu ergänzen.

Literatur

- ANTHROPIC (2025): „Get Started. What is the Model Context Protocol (MCP)?“ <https://modelcontextprotocol.io> (26.11.2025).
- BALDASSARE, D. (2025): „Wissens-Turbo für Chatbots – So wird Technische Dokumentation zum Knowledge Graph für KI.“ Vortrag, gehalten auf der tekomp-Jahrestagung, Stuttgart. Stuttgart.
- DIN 82079 (2021): *Erstellung von Nutzungsinformationen (Gebrauchsanleitungen) für Produkte. Teil 1: Grundsätze und allgemeine Anforderungen*. Berlin: Beuth.
- DREWER, P. / ZIEGLER, W. (2014): *Technische Dokumentation: Übersetzungsgerechte Texterstellung und Content-Management* (2., überarbeitete und aktualisierte Auflage). Würzburg: Vogel.
- FERRARI, D. / FLEISCHMANN, K. (2025): „Terminologie – unentbehrliche Ressource im KI-Zeitalter.“ Vortrag, gehalten auf der tekomp-Jahrestagung, Stuttgart. Stuttgart.
- FLEISCHMANN, K. / LANG, C. (2025): „Terminologie in der KI. Wie mit Terminologie der Output von LLMs und GenAI optimiert werden kann.“ *Terminologie in der KI – KI in der Terminologie. Akten des Symposions. Worms, 27.–29. März 2025*. Hrsg. P. Drewer / F. Mayer / D. Pulitano. München/Karlsruhe/Bern: Deutscher Terminologie-Tag e. V. 83–95.

- GAO, Y. / XIONG, Y. / GAO, X. / JIA, K. / PAN, J. et al. (2024): „Retrieval-Augmented Generation for Large Language Models: A Survey.“ v5, 27 März 2024. 1–21. <https://arxiv.org/abs/2312.10997> (26.11.2025).
- GILLHAUSSEN, T. VON / BAUR, P. / KRECKWITZ, S. (2024): „Wie KI die Terminologiearbeit verändert.“ *eDITion* 20.1: 31–37.
- HAMM, J. (2025): „Terminologische Konsistenz und generative KI – ein *Perfect Match*? Produktiver Einsatz von Sprachmodellen im Terminologiemanagement und beim Post-Editing.“ *Terminologie in der KI – KI in der Terminologie. Akten des Symposiums. Worms, 27.–29. März 2025*. Hrsg. P. Drewer / F. Mayer / D. Pulitano. München/Karlsruhe/Bern: Deutscher Terminologie-Tag e.V. 151–163.
- HAZLEDINE, T. (2024): „Alternatives to Cosine Similarity.“ <https://tomhazledine.com/cosine-similarity-alternatives/> (26.11.2025).
- HUMMER, J. (2023): „M-RAG: Retrieval Augmented Generation Goes Global. Revolutionizing Global Customer Support: The Power of Multilingual Retrieval Augmented Generation (M-RAG).“ <https://coreon.com/2023/12/05/m-rag-goes-global/> (26.11.2025).
- ISO 30042 (2019): *Management of Terminology Resources – TermBase eXchange (TBX)*. Genf: ISO.
- JURAFSKY, D. / MARTIN, J. H. (2025): *Speech and Language Processing. An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition* (3rd edition, draft August 24, 2025). <https://web.stanford.edu/~jurafsky/slp3/> (26.11.2025).
- KRÜGER, R. (2025): „Sprachbezogene Künstliche Intelligenz in Form von Large Language Models. Funktionsweise – fachkommunikatives Anwendungsspektrum – Kompetenzanforderungen.“ *Terminologie in der KI – KI in der Terminologie. Akten des Symposiums. Worms, 27.–29. März 2025*. Hrsg. P. Drewer / F. Mayer / D. Pulitano. München/Karlsruhe/Bern: Deutscher Terminologie-Tag e.V. 11–22.
- LACKNER, A. / VEGA-WILSON, A. / LANG, C. (2025a): „Terminology Augmented Generation: A Systematic Review of Terminology Formats for In-Context Learning in LLMs.“ *Proceedings of the 4rd International Conference on Multilingual Digital Terminology Today (MDTT 2025)*. Thessaloniki. Hrsg. F. Vezzani / G. M. Di Nunzio / E. Loupaki / G. Meditskos / M. Papoutsoglou. o. S.

- LACKNER, A. / VEGA-WILSON, A. / LANG, C. (2025b): „An Evaluation of Terminology-Augmented Generation (TAG) and Various Terminology Formats for the Translation Use Case.“ *Journal of Digital Terminology and Lexicography* 1.2: 97–104.
- LANG, C. / SUCHOWOLEC, K. (2020): „Wissensmanagement in der Praxis: Welchen Beitrag leistet deskriptive Terminologearbeit?“ *Interdependenzen und Innovationen in Translation und Fachkommunikation*. Hrsg. B. Ahrens / M. Beaton-Thome / M. Krein-Kühle / R. Krüger / L. Link / U. Wienen. Berlin: Frank & Timme. 17–44.
- LANG, C. / SUCHOWOLEC, K. / JOCHUM, V. (eingereicht): „The More the Better? Terminological Data as Knowledge Base for a RAG-based Question Answering.“
- LANGCHAIN (o.J.): „Rlm/rag-prompt.“ <https://smith.langchain.com/hub/rlm/rag-prompt> (26.11.2025).
- LEWIS, P. / PEREZ, E. / PIKTUS, A. / PETRONI, F. / KARPUKHIN, V. et al. (2021): „Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks.“ v4, 12 April 2021. 1–19. <https://arxiv.org/abs/2005.11401> (26.11.2025).
- LI, Z. / LI, C. / ZHANG, M. / MEI, Q. / BENDERSKY, M. (2024): „Retrieval Augmented Generation or Long-Context LLMs? A Comprehensive Study and Hybrid Approach.“ *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing: Industry Track*. Eds. Y. Al-Onaizan / M. Bansal / Y.-N. Chen. Association for Computational Linguistics. 881–893. <https://aclanthology.org/2024.emnlp-industry.66.pdf>. (26.11.2025).
- NESBIGALL, J. (2025): „Glossarerstellung für MT und KI. Warum viel nicht immer viel hilft.“ *Terminologie in der KI – KI in der Terminologie. Akten des Symposions. Worms, 27.–29. März 2025*. Hrsg. P. Drewer / F. Mayer / D. Pulitano. München/Karlsruhe/Bern: Deutscher Terminologie-Tag e.V. 51–60.
- NUNZIO, G. M. DI (2025): „Terminology-Augmented Generation (TAG): Foundations, Use Cases, and Evaluation Paths.“ *Journal of Digital Terminology and Lexicography* 1.1: 97–104.
- REIMERINK, A. / ARAÚZ, P. L. / MAGAÑA REDONDO, P. J. (2010): „EcoLexicon: An Environmental TKB.“ *Proceedings of the Seventh International Conference on Language Resources and Evaluation (LREC'10), 19–21 May 2010, Valetta, Malta*. Hrsg. N. Calzolari / K. Choukri / B. Maegaard / J. Mariani / J. Odijk et al. Paris: LRA. 2323–2329.
- REINEKE, D. (2025): „Glossare für die neuronale maschinelle Übersetzung.“ *eDITion* 21.1: 19–25.

- SCHULHOFF, S. / ILIE, M. / BALEPUR, N. / KAHADZE, K. / LIU, A. et al. (2024): „*The Prompt Report: A Systematic Survey of Prompt Engineering Techniques*.“ v6 26. Februar 2025. 1–80. <https://arxiv.org/pdf/2406.06608>. (26.11.2025).
- STOCK, W. G. / STOCK, M. (2015): *Handbook of Information Science*. Berlin/Boston: De Gruyter Saur.
- VASWANI, A. / SHAZEER, N. / PARMAR, N. / USZKOREIT, J. / JONES, L. et al. (2017): „Attention is All You Need.“ *Advances in Neural Information Processing Systems (NIPS 2017)*. Eds. I. Guyon / U. von Luxburg / S. Bengio / R. Fergus / S. V. N. Vishwanathan / R. Garnett. 1–11. https://proceedings.neurips.cc/paper_files/paper/2017/file/3f5ee243547dee91fbd053c1c4a845aa-Paper.pdf. (26.11.2025).
- XU, P. / PING, W. / WU, X. / MCAFEE, L. / ZHU, C. et al. (2024): „Retrieval Meets Long Context Large Language Models.“ 1–16. <https://arxiv.org/abs/2310.03025> (26.11.2025).
- YOU, K. (2025): „Semantics at an Angle: When Cosine Similarity Works Until It Doesn’t.“ v2 20. Mai 2025. o. S. <https://arxiv.org/pdf/2504.16318> (26.11.2025).
- ZHAO, W. X. / ZHOU, K. / LI, J. / TANG, T. / WANG, X. et al. (2023): „A Survey of Large Language Models.“ v16 11 März 2025. 1–144. <https://arxiv.org/abs/2303.18223> (26.11.2025).
- ZHONG, Y. / WEN, Y. / GUO, J. / KAFI, M. / HUANG, H. et al. (2025): „Web Intellectual Property at Risk: Preventing Unauthorized Real-Time Retrieval by Large Language Models.“ *Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing*. Eds. C. Christodoulopoulos / T. Chakraborty / C. Rose / V. Peng. 17223–17236. Suzhou, China: Association for Computational Linguistics. <https://aclanthology.org/2025.emnlp-main.870.pdf>. (26.11.2025).

FACHKOMMUNIKATIONSFORSCHUNG
MIT UND ZU KI

Rechtstranslatorische Verfahren der Künstlichen Intelligenz

Frame-semantische Überlegungen zur Übertragung kulturspezifischer Rechtsbegriffe bei der KI-generierten Urkundenübersetzung

Abstract

Despite the growing importance of Artificial Intelligence across the various branches of linguistics, its benefits for legal translation have not yet been sufficiently explored. However, recent research has revealed deficiencies, specifically in terms of terminological congruence, semantics, and adequate communication of knowledge. This paper proposes a case-specific study of AI-generated translations of a French court ruling into the German legal language, taking into account the requirements for court-certified translations and the system-dependent legal knowledge of expert recipients. Drawing on frame-semantic approaches to cross-systemic communication in law, this study analyzes the linguistic signs chosen by different AI engines in order to provide access to highly culture-specific legal terms, focusing on the potential knowledge elements evoked by legal experts within the source and target systems and thus providing observations on the practical benefits of AI for certified legal translators.

Keywords Rechtsterminologie, Frame-Semantik, Rechtsübersetzung, Übersetzung und Künstliche Intelligenz, Rechtslinguistik

1 Einleitung

Obwohl sich die Künstliche Intelligenz mittlerweile im Übersetzungssektor etabliert hat, ist die empirische Erforschung ihres Nutzens im Hinblick auf das juristische Fachgebiet nicht auf der Höhe der Zeit. So erkennen Bajčić/Golenko (2024) einen erheblichen Bedarf an neuen qualitativen rechtstranslatologischen Studien, und zwar mit Fokus auf Herausforderungen im Bereich der Wissensvermittlung und intratextuellen Kohärenz. Auch wenn der Mehrwert der KI nach Briva-Iglesias (2025) für diesen Tätigkeitsbereich durchaus festzustellen ist und eine immer größere Bedeutung einnimmt, sind bei KI-generierten Rechtsübersetzungen laut Wiesmann (in Vorbereitung) weiterhin zahlreiche Sinnfehler anzutreffen. Vor allem im – stets kontrastiven – rechtssystemübergreifenden Kontext kommen diesbezüglich spezifische Herausforderungen zum Vorschein, die bei der Übersetzung innerhalb mehrsprachiger Rechtsordnungen nicht auffallen (Monjean-Decaudin 2025).

Vor diesem Hintergrund bietet die vorliegende Untersuchung frame-semantische, d. h. kognitive und wissensorientierte, Einblicke in KI-generierte Rechtstranslate, und zwar mit Blick auf besonders kulturspezifische Rechtskonzepte. Ziel ist es, das beim Rezipienten des Zieltextes hervorgerufene Wissen zu beschreiben, das durch die Wahrnehmung der von KI-Modellen gewählten sprachlichen Zeichen bzw. Übersetzungsverfahren bedingt ist. Auf dieser Grundlage lassen sich fallbezogene Überlegungen zum Mehrwert der KI bei der Zugänglichmachung rechtssystemgebundenen Wissens formulieren.

2 Juristische Fachübersetzung als Kommunikation von Wissen

In der Fachsprachenforschung nimmt die Rechtssprache oft eine gesonderte Stellung ein, was zuvorderst auf ihren abstrakten Charakter, eine notwendige Interpretationsoffenheit, ihre unaufhaltbare Dynamizität, die Normativität der juristischen Semantik und nicht zuletzt ihre extreme Systemgebundenheit zurückzuführen ist (Nazarov 2025: 19–70). Die Übersetzung von Rechts-

texten ist somit von angepassten Methodologien geprägt, die allesamt einen Rechtsvergleich zwischen dem Verstehens- und dem Transferprozess verorten (Pommer 2006; Bocquet 2008; Monjean-Decaudin 2022). Dabei wird in erster Linie auf die Etablierung einer funktionalen Äquivalenz hingewirkt, die darin besteht, im Zielrechtssystem hochgradig ähnliche Aufgaben erfüllende Konzepte zu eruieren, die allerdings aufgrund einer möglichen Verzerrung der Ausgangsrechtsordnung mit Gefahren in Verbindung gebracht wird (Harvey 2002: 42–43). Es entstehen somit parallele Übersetzungsverfahren, die teilweise als Ersatzlösungen bezeichnet werden (de Groot 1999: 27), darunter eine deskriptive oder formelle Übertragung sowie die Nichtübersetzung, die Bildung von Neologismen und auch die Paraphrasierung, wie etwa von Stolze (2014: 273) typologisiert.

Da die mögliche Etablierung von Eins-zu-eins-Entsprechungen zwischen verschiedenen Rechtsordnungen grundsätzlich abgelehnt wird (Pommer 2006: 64–65), kann bei der systemübergreifenden Übersetzung keine absolute semantische Konvergenz zwischen Ziel- und Ausgangstext erwartet werden. Vielmehr ist dem von Engberg (2013: 12) postulierten Ziel beizupflichten, das darin besteht, den Lesern der Zielrechtsordnung das Wissen aus dem Ausgangstext durch Kommunikation zugänglich zu machen, was vor allem durch eine frame-semantische Betrachtung beschreibbar wird, die das von Rezipienten hervorgerufene Wissen fokussiert. Bei der beglaubigten Übersetzung von rechtssprachlichen Urkunden, wie Gerichtsurteilen, die von beeidigten Übersetzern angefertigt werden, kommt dies besonders zum Vorschein. Diese ermöglicht es, die bundesdeutsche Richterin als prototypische sowie auch endgültige Rezipientin und gleichzeitig Expertin des juristischen Fachgebiets heranzuziehen, die eine beglaubigte Übersetzung des französischen Urteils im Zuge eines deutschen Gerichtsverfahrens rezipiert, auslegt und ggf. anwendet. Für die Untersuchung wird genau dieser Rahmen gewählt, da dabei ein besonderes Augenmerk auf die Übertragung stark kulturspezifischer Rechtskonzepte gelegt wird.

3 Eine frame-semantische Untersuchung KI-generierter Rechtstranslate

Im Rahmen der vorliegenden Untersuchung wird ein kleines Korpus von Translaten desselben Rechtstextes zusammengestellt, die von sechs verschiedenen KI-Modellen generiert wurden: ChatGPT, Claude Sonnet 4, Copilot Think Deeper, Gemini Flash 2.5, Mistral AI – Le Chat und Perplexity. Beim Ausgangstext handelt es sich um das bereits in Nazarov (2025) terminologisch untersuchte Gerichtsurteil *arrêt n° 548 FS-B – pourvoi n° B 20-21.103*, das von der *Cour de Cassation*, dem obersten Gericht der ordentlichen Gerichtsbarkeit Frankreichs, erlassen wurde und sich mit einer Entscheidung des vorinstanzlichen Berufungsgerichts, der *cour d’appel*, im Rahmen einer Revision befasst, welches wiederum über die Auflösung eines Arbeitsverhältnisses entschied.

Folgender Übersetzungsauftrag (*Prompt*) wurde den verwendeten KI-Modellen erteilt, der den hier beschriebenen Rahmen deutlich macht:

Der vorliegende zu übersetzende Ausgangstext stellt einen Rechtstext dar, der der Rechtsordnung Frankreichs entstammt. Dieser soll in die deutsche Rechtssprache der Bundesrepublik Deutschland übersetzt werden. Bitte fertige eine juristische Übersetzung nach den Maßgaben einer beglaubigten Urkundenübersetzung an, sodass der Zieltext – im Falle einer Beglaubigung durch einen beeidigten Übersetzer – in einem deutschen Gerichtsverfahren vorgelegt und angewandt werden könnte. Bitte übersetze Satz für Satz.

Aus frame-semantischer Perspektive soll anhand ausgewählter stark kulturspezifischer Begriffe¹ beschrieben werden, welches Wissen bei Rezipienten der Zielrechtsordnung kognitiv hervorgerufen wird, wenn die von den jeweiligen KI-Tools gewählten sprachlichen Zeichen innerhalb des Translats wahrgenommen werden, und inwiefern dieses die im Einklang mit dem ausgangs-

.....

- 1 Im vorliegenden Beitrag werden Rechtstermini, die bereits in Nazarov (2025) aus frame-semantischer Sicht akribisch untersucht wurden, nur in prägnanter Form vorgestellt. Für detaillierte Ausführungen und juristische Quellen sei auf dieses Werk verwiesen.

sprachlichen Verwendungskontext aktivierten relevanten Frame-Elemente widerspiegelt.

3.1 *rupture conventionnelle*

Für die frame-ontologische Beschreibung des innerhalb der Ausgangsrechtsordnung evozierten Wissensrahmens rund um den Begriff *rupture conventionnelle* wird relevant, dass diese eine mögliche Form – d. h. ein Subframe – der übergeordneten *rupture du contrat de travail* darstellt und somit parallel zu anderen Optionen wie *résiliation judiciaire* positioniert werden muss. Legal definiert wird das Konzept u. a. in Artikel L. 1237–11 des *Code du travail* als zwischen Arbeitnehmer und Arbeitgeber geschlossene Vereinbarung darüber, ein bestehendes Arbeitsverhältnis einvernehmlich zu beenden. Die Verpflichtungen und Folgen, die in deren Rahmen entstehen, wie etwa eine Abfindungszahlung, sind frame-semantisch erfassbar. Zudem wird auch ein bestimmtes Verfahren auferlegt, das bestimmte *entretiens* vor Vollzug der *rupture* vorsieht, bei denen die Parteien eine schriftliche Vereinbarung ausarbeiten, die einer Prüfung und *homologation* durch die sogenannten *DREETS* unterliegt. Durch diese Auslegung ergibt sich bereits eine einzigartige ontologische Konstellation der hier genannten Konzepte, die als Subframes bzw. eigenständige Frames miteinander verknüpft und einzeln granulierbar sind.

Bei der Übertragung dieses Konzepts wird, wie in Tabelle 1 ersichtlich, tendenziell eine funktionale Äquivalenz bevorzugt, und zwar durch die sprachlichen Zeichen *Aufhebung* bzw. *Auflösung*, die im bundesdeutschen Rechtssystem verankert sind. Ein Aufhebungs- bzw. Auflösungsvertrag bewirkt konkret eine einvernehmliche Beendigung des Arbeitsverhältnisses, die sowohl für Arbeitnehmer als auch Arbeitgeber vorteilhaft sein kann, da Kündigungsschutz und -frist entfallen. Das durch diese Zeichen evozierte Wissen ermöglicht somit den Zugang zum kulturspezifischen französischen Konzept, wobei auffällig abweichende Frame-Elemente und -Relationen (wie z. B. die Trennung vom deutschen Kündigungsschutz und sozialrechtliche Sperrzeit beim deutschen Konzept einerseits und das in *entretiens* und *homologation* bestehende

Verfahren und die Abfindungspflicht des französischen Frames andererseits) im Verwendungskontext keine Verständnisprobleme bereiten.

Tab. 1: KI-Wiedergabe des Begriffs *rupture conventionnelle*

ChatGPT – GPT-4o	• <i>einvernehmliche Aufhebungsvereinbarung</i> • <i>Aufhebungsvereinbarung</i> • <i>Aufhebungsvereinbarung (rupture conventionnelle)</i>
Claude Sonnet 4	• <i>Aufhebungsvergleich</i> • <i>Aufhebungsvereinbarung</i>
Copilot Think Deeper	• <i>Auflösungsvereinbarung</i> • <i>Vereinbarung</i> • <i>Auflösung</i> • <i>einvernehmliche Auflösung</i>
Gemini Flash 2.5	• <i>Aufhebungsvereinbarung</i> • <i>einvernehmliche Auflösung</i>
Mistral AI – Le Chat	• <i>Aufhebungsvereinbarung</i> • <i>Aufhebung</i> • <i>einvernehmliche Beendigung</i>
Perplexity	• <i>Aufhebungsvereinbarung</i> • <i>einvernehmliche Aufhebung</i>

Die KI-Modelle wählen in erster Linie die sprachlichen Zeichen der funktionalen Äquivalenz *Aufhebungs-* bzw. *Auflösungsvereinbarung*. Konstellationen wie *einvernehmliche Aufhebungsvereinbarung*, die auf die Übertragung des Adjektivs bestehen, finden sich bei allen außer bei den von Claude und Mistral AI generierten Translaten ebenfalls, erscheinen allerdings redundant, da die Einvernehmlichkeit bereits durch das bundesdeutsche Rechtskonzept evoziert wird und das Adjektiv *conventionnelle* zur Konstruktion des französischen Rechtsterminus gehört. Die Verwendung *einvernehmliche Beendigung* bei Mistral AI kann sowohl als formelle als auch als deskriptive Übertragung gesehen werden, erscheint innerhalb des generierten Zieltextes allerdings insofern problematisch, als an anderen Stellen die funktionalen Alternativen *Aufhebung* und *Aufhebungsvereinbarung* gewählt werden, was zu einer intratextuellen terminologischen Inkohärenz führt. Dies ist allerdings streng genommen bei

allen verwendeten Tools festzustellen, da jeweils mindestens zwei Varianten angeführt werden. Auffällig ist bei ChatGPT, dass es eine Kombination aus funktionaler Äquivalenz und Nichtübersetzen – *Aufhebungsvereinbarung* (*rupture conventionnelle*) – anbietet. Ein solches Vorgehen ist bei stark kulturspezifischen Konzepten, Eigennamen und Institutionenbezeichnungen im Kontext der beglaubigten Übersetzung gängig und teilweise notwendig. Fraglich ist allerdings, ob der behandelte Begriff tatsächlich so einzigartig und problembehaftet ist, dass die französische Benennung zwingend angeführt werden muss, und warum dieses Übersetzungsverfahren nur einmal vorkommt, und zwar erst nachdem im selben Translat bereits Varianten ohne das Nichtübersetzen verwendet werden. Der einzige Vorschlag, der sich durchaus als problematisch erweisen kann, ist der von Claude im Kontext der Abfindung verwendete *Aufhebungsvergleich*, da die Aufhebungsvereinbarung vom Konzept des Vergleichs zu trennen ist (siehe Kapitel 3.3).

3.2 *Cour de cassation, chambre sociale*

Die nächsten untersuchten Konzepte sind die institutionellen Benennungen *Cour de cassation* und *chambre sociale*, die hier in enger Verbindung zueinander betrachtet werden müssen. Als *juridiction suprême de l'ordre judiciaire* untersucht das Gericht die Einhaltung der *règle de droit* vorinstanzlicher Rechtsprechungsorgane, in erster Linie „en se prononçant sur les pourvois dirigés contre les arrêts rendus par les cours d'appel“ (Perrot et al. 2022: 170–171). Neben diesen *missions contentieuses* verfolgt die *Cour de cassation* auch sogenannte *missions consultatives* und entscheidet zudem über die Vorlage der verfassungsrechtlichen *question prioritaire de constitutionnalité* (QPC) beim *Conseil constitutionnel*. Die *chambre sociale*, die für arbeitsrechtliche Angelegenheiten zuständig ist und sich parallel zu anderen *chambres* verorten lässt, erscheint als Frame-Element und somit als eigenes Subframe des Wissenssegments, das durch die *Cour de cassation* evoziert wird. Dieses schließt auch weitere Bestandteile ein wie den *greffier* und die verschiedenen Arten von *conseiller*, die auch im Ausgangstext an verschiedenen Stellen aufgeführt werden.

Die Übertragung dieses französischen Frames in die bundesdeutsche Rechtssprache stellt insofern Herausforderungen dar, als das *Bundesarbeitsgericht* für diesen Kontext hochgradig ähnliche Aufgaben erfüllt und somit als funktionales Äquivalent für die *chambre sociale de la Cour de cassation* infrage kommt. Für den Begriff der *Cour de cassation* könnte in diesem Zusammenhang wiederum der *Bundesgerichtshof* herangezogen werden. In beiden Fällen werden die mit der funktionalen Äquivalenz assoziierten Gefahren erkennbar, die hier u. a. im föderalistischen System und den sehr unterschiedlichen Gerichtsbarkeiten bestehen. Alle KI-Tools außer Perplexity bilden Ausdrücke mit *Kassation*, und zwar entweder *Kassationshof* oder *Kassationsgerichtshof* (siehe Tabelle 2).

Tab. 2: KI-Wiedergabe der Begriffe *Cour de cassation* und *chambre sociale*

ChatGPT – GPT-4o	• <i>COUR DE CASSATION – KASSATIONSGERICHTSHOF</i> • <i>KASSATIONSGERICHTSHOF</i>
	• <i>Sozialkammer</i>
Claude Sonnet 4	• <i>Kassationshof</i>
	• <i>Sozialsenat</i>
Copilot Think Deeper	• <i>Cour de Cassation</i> • <i>Kassationshof</i>
	• <i>Sozialkammer</i> • <i>chambre sociale</i>
Gemini Flash 2.5	• <i>KASSATIONSHOF</i>
	• <i>Sozialkammer</i>
Mistral AI – Le Chat	• <i>KASATIONSHOF</i> • <i>Kassationshof</i>
	• <i>Sozialkammer</i>
Perplexity	• <i>Oberster Gerichtshof</i>
	• <i>Sozialkammer</i>

Einerseits kann dies als formelle Übertragung und Neologismus betrachtet werden, da es keine solche Institution in Deutschland gibt, oder auch als Über-

nahme der existierenden schweizerisch-deutschen Bezeichnung, was allerdings nicht mehr mit dem vorgegebenen bundesdeutschen Kontext übereinstimmt. Andererseits kann das Konzept der Kassation in Deutschland auch als funktional gelten, da es im mittelalterlichen Recht als Terminus durchaus etabliert war und aus historischer Sicht Wissen evoziert werden kann. Wie auch beim vorherigen Begriff (siehe Kapitel 3.1) führt ChatGPT, diesmal bei der ersten Verwendung, zusätzlich die ausgangssprachliche Benennung an und hält sich dabei an die für beglaubigte Übersetzungen übliche Übertragung von Institutionenbezeichnungen. Copilot scheint den Begriff ebenfalls als Institution zu erkennen und behält diesen in der französischen Sprache bei, ohne somit einen Zugang zum Wissenssegment aus der Perspektive der deutschen Rechtssprache zu schaffen. Die (inkohärente) Verwendung der orthografisch nicht korrekten Schreibweise *Kasationshof* bei Mistral AI lässt sich nur schwer nachvollziehen. Perplexity ist wiederum das einzige Tool, das auf die Verwendung der Kassation verzichtet und konsequent eine deskriptive Äquivalenz in Form von *Oberster Gerichtshof* wählt, wobei Wissen über die hierarchische Stellung des Rechtsprechungsorgans vermittelt und die Aufgabe der Prüfung vorinstanzlicher Gerichte suggeriert werden.

Was die Übertragung der arbeitsrechtlichen *chambre sociale* angeht, wird tendenziell die formelle Äquivalenz *Sozialkammer* gewählt, die auch als Neologismus in Erscheinung tritt, da es eine solche in der bundesdeutschen Rechtsordnung nicht gibt. Dabei suggeriert dieses verwendete sprachliche Zeichen eine Anbindung an das Sozialrecht, das in Deutschland vom Arbeitsrecht zu trennen ist, wo sich diese beiden Rechtsgebiete in jeweils eigenen Gerichtsbarkeiten manifestieren. Der französische *droit social* beinhaltet wiederum das Arbeitsrecht, *droit du travail*, was durch die deutsche Bezeichnung nicht zugänglich wird. Copilot verwendet zum Teil dasselbe Übersetzungsverfahren, wählt an einer späteren Stelle jedoch lediglich das Nichtübersetzen (*chambre sociale*), während ChatGPT dieses Organ nicht in der ausgangssprachlichen Version wiedergibt, obwohl diese Tendenz bei diesem Modell sehr deutlich wird. Auffällig ist, dass Claude als einziges Tool den Neologismus *Sozialsenat* bildet und diesen kohärent verwendet, was sich damit erklären lässt, dass weder der Bundesgerichtshof noch das Bundesarbeitsgericht über *Kammern* verfügen, wie es bei niedriger gestellten Gerichten der Fall ist, sondern über

Senate. Die Wahl dieses sprachlichen Zeichens evoziert somit adäquates Wissen über die Stellung und Funktion des übertragenen Organs, allerdings bildet die Kollokation mit *Sozial-* weiterhin ein Wissenssegment, das das Arbeitsrecht nicht widerspiegelt.

3.3 *convocation en bureau de conciliation*

Die Konstruktion *convocation en bureau de conciliation* entspricht im Kontext des gewählten Urteils dem Arbeitsprozessrecht, da die Bemühung um die Erreichung einer *conciliation* vorgeschrieben ist, weshalb die Parteien den ersten Schritt im *bureau de conciliation* am erstinstanzlichen Arbeitsgericht *conseil de prud'hommes* vollziehen, bevor ggf. das streitige Verfahren im *bureau de jugement* beginnt. Alle KI-Tools außer Copilot wählen strikt Varianten einer funktionalen Äquivalenz wie etwa *Ladung zum Güutetermin* und verzichten dabei auf eine formelle Übertragung bzw. die Angabe der ausgangssprachlichen Bezeichnung (siehe Tabelle 3). Laut § 54 ArbGG beginnt die mündliche Verhandlung am Arbeitsgericht zwingend mit einem Güutetermin, und zwar mit dem Ziel, eine güutliche Streitbeilegung in Form eines Vergleichs zu erreichen, was im Falle eines Scheiterns den Anschluss eines Kammertermins bzw. einer Kammerverhandlung nach sich zieht. Auch wenn die einzelnen Elemente der prozessrechtlichen Konzepte wie die Zuordnung zu einem *bureau* diese semantisch inkongruent machen, besteht eine funktionale Vergleichsgrundlage, die die KI-Tools scheinbar als fruchtbar werten. Überraschend ist, dass Copilot, trotz Anweisung im Übersetzungsauftrag, Satz für Satz zu übertragen, sich für eine Zusammenfassung des gesamten Absatzes entscheidet und dabei das hier behandelte Konzept in keiner Weise wiedergibt. Die Güteverhandlung wird somit beim Rezipienten des von ihm generierten Translats nicht kognitiv-epistemisch hervorgerufen.

Tab. 3: KI-Wiedergabe des Begriffs *convocation en bureau de conciliation*

ChatGPT – GPT-4o	• <i>Ladung zur Güteverhandlung</i>
Claude Sonnet 4	• <i>Ladung zum Güteverfahren</i>

Copilot Think Deeper	• -
Gemini Flash 2.5	• <i>Ladung zur Gütestelle</i>
Mistral AI – Le Chat	• <i>Ladung zur Gütestelle</i>
Perplexity	• <i>Ladung zum Güetermin</i>

3.4 *Code de procédure civile, Code du travail*

Die Begriffe *Code de procédure civile* und *Code du travail* werden in einem Unterkapitel gemeinsam behandelt, da sie derselben Kategorie zuzuordnen sind: Bezeichnungen von Gesetzen bzw., in diesem Fall, Gesetzeswerken. Der *Code de procédure civile* besteht aus sechs *Livres* und regelt neben dem *Code de l'organisation judiciaire* sowie dem *Code des procédures civiles d'exécution* den französischen Zivilprozess. Auch wenn die bundesdeutsche *Zivilprozessordnung*, das funktionale Äquivalent, aus elf *Büchern* besteht, deren Inhalte nicht deckungsgleich mit dem französischen System sind, verwenden ausnahmslos alle KI-Modelle dieses Übersetzungsverfahren, um den französischen Terminus zugänglich zu machen (siehe Tabelle 4). Kein Modell hält es für notwendig, durch formelle Übertragung entstehende Neologismen wie *Zivilprozessgesetzbuch* o.Ä. zu bilden. Eine Besonderheit ist erneut bei ChatGPT zu beobachten, das dieses funktionale Äquivalent durch das Adjektiv *französisch* lexikalisch erweitert und mit der ausgangssprachlichen Benennung verbindet, sodass einerseits zielrechtssystemgebundenes Wissen evoziert wird, andererseits allerdings auch eine Lokalisierung in der Ausgangsrechtsordnung gewährleistet ist.

Tab. 4: KI-Wiedergabe der Begriffe *Code de procédure civile* und *Code du travail*

ChatGPT – GPT-4o	• <i>französische Zivilprozessordnung</i> (<i>Code de procédure civile</i>)
	• <i>französisches Arbeitsgesetzbuch</i> (<i>Code du travail</i>)
	• <i>Arbeitsgesetzbuch</i>

Claude Sonnet 4	• <i>Zivilprozessordnung</i>
	• <i>Arbeitsgesetzbuch</i>
Copilot Think Deeper	• <i>Zivilprozessordnung</i>
	• <i>Code du travail</i>
Gemini Flash 2.5	• <i>Zivilprozessordnung</i>
	• <i>Arbeitsgesetzbuch</i>
Mistral AI – Le Chat	• <i>Zivilprozessordnung</i>
	• <i>Arbeitsgesetzbuch</i>
Perplexity	• <i>Zivilprozessordnung</i>
	• <i>Arbeitsgesetzbuch</i>

Was den zweiten Begriff, den *Code du travail*, betrifft, so muss zwischen den gewählten Vergleichsländern eine besondere Diskrepanz beachtet werden, weil es sich um ein gesetzgeberisches Grundlagenwerk für das französische Arbeitsrecht handelt, wofür es in Deutschland kein funktionales Gegenstück gibt. Der *Code du travail* teilt sich in verschiedene Arten von Paragraphen auf, die jeweils durch die Buchstaben L, R und D voneinander getrennt werden. Ergänzt wird dieser durch andere das Arbeitsrecht prägende Werke wie den *Code de la Sécurité sociale* und den *Code de la mutualité*. Bis auf Copilot, das die ausgangssprachlichen Zeichen unverändert wiedergibt, ohne zusätzlichen Zugang zu ermöglichen, verwenden alle Modelle *Arbeitsgesetzbuch* (siehe Tabelle 4). Da es ein solches in der bundesdeutschen Rechtsordnung nicht gibt, weil das deutsche Arbeitsrecht durch zahlreiche einzelne Quellen wie das Kündigungsschutz- und das Arbeitszeitgesetz geregelt wird, kann diese Übersetzungstechnik als formelle Übertragung und Neologismus betrachtet werden. Gleichzeitig handelt es sich um ein deskriptives Verfahren, da es in Anlehnung an ähnliche Konstruktionen wie Strafgesetzbuch, das aus den Elementen Strafrecht und Gesetzbuch besteht, Wissen darüber vermittelt, dass es sich um ein gesetzgeberisches Grundlagenwerk für das Arbeitsrecht handelt. ChatGPT kombiniert dieses bei der ersten Verwendung analog zur Übertragung des *Code de procédure civile* mit der ausgangssprachlichen Benennung und lokalisiert es durch das Adjektiv *französisch* im Ausgangsrechtssystem.

4 Fazit

Bei dieser kognitiven und wissensorientierten Betrachtung fällt auf, dass ChatGPT durch die Verbindung von funktionaler Äquivalenz und Nichtübersetzen ein der beglaubigten Urkundenübersetzung stark entsprechendes Translat erstellt und so Institutionen wie die *Cour de cassation* nicht nur verständlich, sondern auch auffindbar macht, was allerdings noch nicht systematisch erfolgt. Copilot weist bei stark kulturspezifischen Begriffen eine Tendenz zur reinen Nichtübersetzung auf, was eine Wissensvermittlung unmöglich macht, und erzeugt durch das Auslassen ganzer Sätze zudem wesentliche Wissenslücken. Darüber hinaus sind in einem Fall ein Rechtschreibfehler bei Mistral AI und eine terminologische Sinnentstellung bei Claude Sonnet festzustellen.

Die gewählten KI-Tools weisen im Kontext der Urkundenübersetzung allesamt dort eine Tendenz zur funktionalen Äquivalenz auf, wo diese keine für dieses Übersetzungsverfahren typischen Gefahren birgt, d. h. kein durch die Etablierung im Zielrechtssystem zu stark abweichendes Wissen evoziert. So werden die Termini *rupture conventionnelle* und *Code de procédure civile* problemlos funktional übertragen, während *Cour de cassation* nicht etwa durch Bezeichnungen wie *Bundesgerichtshof* verzerrt wird. Auch Neologismen wie *Arbeitsgesetzbuch*, die sich an bereits in der deutschen Rechtsordnung bestehende Konstruktionen anlehnen, scheinen für die KI insgesamt keine Herausforderung darzustellen. Andererseits deutet die Wiedergabe der *chambre sociale* mit *Sozialkammer* bzw. *-senat* darauf hin, dass auf zusätzliche Schritte zur Wissenslückenschließung, die etwa in Fußnoten erfolgen, verzichtet wird. Letztendlich lässt die terminologische Kohärenz angesichts der Synonymieproblematik im Fachgebiet Recht bei allen KI-Modellen zu wünschen übrig.

Die vorliegende Untersuchung der Wiedergabe ausgewählter Rechtskonzepte in KI-generierten Translaten erhebt nicht den Anspruch auf Allgemeingültigkeit, sondern versteht sich als fallbezogene Betrachtung. Die gewählte Vorgehensweise besteht nämlich darin, die von der KI verwendeten sprachlichen Zeichen mit etablierten rechtstranslatorischen Verfahren in Verbindung zu bringen und frame-semantic zu erfassen, welche Wissensrahmen bei der Zieltextleserin, einer bundesdeutschen RichterIn, zwecks Zugänglichmachung systemgebundenen Wissens evoziert werden. Dies ermöglicht erste

frame-semantische Überlegungen zum KI-generierten Output im Rahmen einer beglaubigten Rechtsübersetzung und zu seinem Nutzen für die Alltags-tätigkeit eines beeidigten Übersetzers.

Literatur

- BAJČIĆ, MARTINA / GOLENKO, DEJANA (2024): „Applying Large Language Models in Legal Translation: The State-of-the-Art.“ *International Journal of Language & Law* 13: 171–196.
- BOCQUET, CLAUDE (2008): *La traduction juridique: Fondement et méthode*. Brüssel: De Boeck.
- BRIVA-IGLESIAS, VICENT (2025): „Are AI Agents the New Machine Translation Frontier? Challenges and Opportunities of Single- and Multi-Agent Systems for Multilingual Digital Communication.“ *Machine Translation Summit XX. Volume 1*. Hrsg. Pierrette Bouillon / Johanna Gerlach / Sabrina Girletti / Lise Volkart / Raphael Rubino / Rico Sennrich / Ana C. Farinha / Marco Gaido / Joke Daems / Dorothy Kenny / Helena Moniz / Sara Szoc. Genf: EAMT. 365–377.
- ENGBERG, JAN (2013): „Comparative Law for Translation: The Key to Successful Mediation between Legal Systems.“ *Legal Translation in Context: Professional Issues and Prospects*. Hrsg. Anabel Borja Albi / Fernando Prieto Ramos. Bern: Lang. 9–25.
- DE GROOT, GERARD-RENÉ (1999): „Das Übersetzen juristischer Terminologie.“ *Recht und Übersetzen*. Hrsg. Gerard-René de Groot / Reiner Schulze. Baden-Baden: Nomos. 11–46.
- HARVEY, MALCOLM (2002): „Traduire l'intraduisible: Stratégies d'équivalence dans la traduction juridique.“ *Revue de l'Institut des langues et cultures d'Europe, Amérique, Afrique, Asie et Australie* 3: 39–49.
- MONJEAN-DECAUDIN, SYLVIE (2022): *Traité de juritraductologie: Épistémologie et méthodologie de la traduction juridique*. Villeneuve d'Ascq: Presses universitaires du Septentrion.
- MONJEAN-DECAUDIN, SYLVIE (2025): „Les enjeux de l'intelligence artificielle pour la traduction juridique au sein de l'Union européenne: Observations générales du point de vue de la juritraductologie.“ *International Journal of Language & Law* 14: 194–203.

- NAZAROV, WALDEMAR (2025): *Frame-basierte Rechtsübersetzung: Frame-Semantik als ontologisches und rechtstranslatorisches Analyseinstrument am Beispiel französischer und bundesdeutscher Rechtstermini*. Berlin: Lang.
- PERROT, ROGER / BEIGNIER, BERNARD / MINIATO, LIONEL (2022): *Institutions juridiques*. Paris La Défense: LGDJ, Lextenso.
- POMMER, SIEGLINDE (2006): *Rechtsübersetzung und Rechtsvergleichung: Translatologische Fragen zur Interdisziplinarität*. Frankfurt am Main: Lang.
- STOLZE, RADEGUNDIS (2014): *Praxishandbuch Urkundenübersetzung: Fertigkeiten, Terminologie, Rechtssprache*. Tübingen: Stauffenburg.
- WIESMANN, EVA (in Vorbereitung): „Intercultural Translation of Legal Texts: A Comparison of MT and AI tools.“ *International Handbook of Legal Language and Communication: From Text to Semiotics*. Hrsg. Anne Wagner. Springer.

Bewertung von Sprachmodellen bei der Übersetzung metaphorischen Denkens in der Wirtschaftssprache

Abstract

This study examines the translation of metaphorical expressions in economic discourse using large language models (LLMs), such as ChatGPT-5, Gemini, and Perplexity. Fifty German metaphors were translated into English and Slovak to compare the models' performance in high- (English) and low-resource (Slovak) languages. Both qualitative and quantitative evaluations were employed to assess metaphorical equivalence and the quality of the translations overall. The results revealed significant differences between the models and language pairs, with proprietary systems demonstrating greater fluency and figurative accuracy. This study deepens our understanding of how LLMs handle figurative meaning and identifies areas for improvement in both high- and low-resource language contexts.

Keywords high-resource language, LLMs, low-resource language, machine translation, metaphor translation

1 Einführung

In den letzten Jahren haben LLMs wie ChatGPT, Gemini und Perplexity den Bereich der maschinellen Übersetzung erheblich vorangebracht und liefern zunehmend flüssige und kontextsensitive Ergebnisse. Trotz dieser Entwicklungen stellt die genaue Übersetzung von Bildsprache, insbesondere von Metaphern, weiterhin eine erhebliche sprachliche und kulturelle Herausforderung

dar. Diese Studie untersucht, inwieweit ausgewählte LLMs in der Lage sind, die metaphorische Bedeutung im wirtschaftlichen Diskurs zu bewahren. Mit Schwerpunkt auf Deutsch als Ausgangssprache und Englisch und Slowakisch als Zielsprachen kombiniert die Forschung qualitative und quantitative Bewertungen, um die Übersetzungsqualität und die bildliche Entsprechung zu beurteilen.

2 Theoretische Ausgangspunkte

2.1 LLMs als Übersetzungstools

Große Sprachmodelle (LLMs) haben die maschinelle Übersetzung grundlegend verändert, da sie flüssige, kontextangemessene und oft überdurchschnittlich präzise Ergebnisse liefern. In ressourcenreichen Sprachpaaren erreichen Modelle wie GPT-4 oder LLaMA dank umfangreicher Trainingsdaten und der Fähigkeit, längere Diskurszusammenhänge zu erfassen, bereits eine Qualität, die mit neuronalen oder kommerziellen Systemen wie Google Translate vergleichbar ist (He et al. 2023; Zhu et al. 2023; Wang et al. 2024). Mithilfe von Kontextlernen und Prompt Engineering lassen sie sich zudem effizient an neue Domänen anpassen (Moslem et al. 2023). In ressourcenarmen Sprachen hingegen bleibt die Leistung begrenzt. Datenknappheit, sprachliche Vielfalt und Verzerrungen führen zu Fehlausrichtungen, Halluzinationen und Qualitätsverlusten (Tars et al. 2022; Tan/Zhu 2024; Tran et al. 2024; Gain et al. 2025; Kalaš 2025a; Kalaš 2025b; Kalaš/Lipták 2025; Song et al. 2025). Um diese Lücken zu verringern, kombinieren aktuelle Ansätze Transferlernen, synthetische Datengenerierung und parameter-effiziente Anpassungstechniken wie LoRA oder Neuronenselektion (Abdelghaffar et al. 2022; Cassano et al. 2023; Liang et al. 2025). Ergänzend tragen strukturierte Prompting-Methoden wie MAPS oder TALENT sowie interaktive, feedbackbasierte Systeme zur Verbesserung der metaphorischen und semantischen Kohärenz bei (He et al. 2023; Gómez et al. 2024; Guo et al. 2024).

2.2 LLMs und idiomatische Sprache

Die Übersetzung bildlicher Sprache erfordert weit mehr als eine wörtliche Wiedergabe, da Redewendungen und Metaphern kulturell geprägt und nicht kompositorisch sind (Li et al. 2023; Castaldo/Monti 2024; Donthi et al. 2024; Zhu et al. 2024). Aktuelle Studien zeigen jedoch, dass LLMs solche Ausdrücke zunehmend angemessen interpretieren und übertragen können. Durch semantisches Denken und Kontextbewusstsein erfassen sie idiomatische Strukturen und leiten die beabsichtigte Bedeutung ab (Li et al. 2023; Qian 2024). Die Nutzung von Wissensdatenbanken wie IdiomKB verbessert dabei die kontextspezifische Genauigkeit (Donthi et al. 2024). Zudem fördern gezieltes Prompt-Design und Paraphrasierungsmethoden wie PASS die stilistische Natürlichkeit und metaphorische Kompetenz (Khoshtab et al. 2024; Gao et al. 2025). Studien belegen, dass fortschrittliche Modelle wie GPT-4 idiomatische Übersetzungen teilweise besser liefern als herkömmliche Systeme, selbst in weniger verbreiteten Sprachen (Ichien et al. 2023; Zhu et al. 2024). Dennoch bleiben Herausforderungen bestehen: Open-Source-Modelle schneiden in ressourcenarmen Sprachen schwächer ab, und stilistische Nuancen oder emotionale Töne gehen häufig verloren. Zudem erfassen gängige Bewertungsmetriken die Feinheiten bildlicher Sprache nur unzureichend. Das macht neue, humanorientierte Bewertungsansätze erforderlich (Li et al. 2023; Donthi et al. 2024; Zhu et al. 2024).

Angesichts dieser Entwicklungen und anhaltender Herausforderungen zielt die vorliegende Studie darauf ab, die Effektivität moderner großer Sprachmodelle bei der Übersetzung metaphorischer Ausdrücke im wirtschaftlichen Diskurs zu bewerten. Dazu werden die drei KI-Systeme ChatGPT, Gemini und Perplexity in zwei Übersetzungsrichtungen (Deutsch–Englisch und Deutsch–Slowakisch) untersucht. Dementsprechend werden in dieser Studie folgende Hypothesen aufgestellt:

H₁: Große Sprachmodelle unterscheiden sich erheblich in ihrer Fähigkeit, metaphorische Bedeutungen zu bewahren und eine hohe Gesamtübersetzungsqualität zu gewährleisten. Dabei werden proprietäre

Modelle (ChatGPT-5 und Gemini) voraussichtlich besser abschneiden als frei zugängliche Systeme (Perplexity).

H₂: Die Übersetzungsleistung wird zwischen ressourcenreichen und ressourcenarmen Sprachen erheblich variieren, wobei für das ressourcenreiche Sprachpaar (Deutsch–Englisch) eine höhere Qualität zu erwarten ist als für das ressourcenarme Sprachpaar (Deutsch–Slowakisch).

H₃: Es wird eine signifikante Wechselwirkung zwischen Modelltyp und Sprachressourcenlevel geben. Das deutet darauf hin, dass bestimmte Modelle je nach Ressourcenstatus der Sprache überproportional besser oder schlechter abschneiden können.

3 Methodologie

In dieser Studie wurden fünfzig metaphorische Ausdrücke aus der Rubrik „Finanzen – Börse und Geldanlage“ des Magazins *Wirtschaftswoche* KI-gestützt übersetzt. Die Sätze wurden mit den Tools ChatGPT-5, Gemini und Perplexity aus dem Deutschen ins Englische und Slowakische übersetzt. Im Anschluss erfolgte eine kombinierte qualitative und quantitative Analyse.

Die qualitative Bewertung folgte dem Rahmen von Wang et al. (2024) und konzentrierte sich auf zwei Hauptdimensionen: metaphorische Äquivalenz und Übersetzungsqualität. Die Äquivalenz wurde als vollständig, teilweise oder nicht gegeben klassifiziert (keine Äquivalenz). Die Übersetzungsqualität wurde anhand der Kriterien Flüssigkeit, Verständlichkeit, Genauigkeit und Authentizität auf einer 5-Punkte-Likert-Skala (1 = sehr schlecht; 5 = ausgezeichnet) bewertet. Bei der quantitativen Analyse wurden deskriptive Metriken eingesetzt, um Leistungsunterschiede zwischen den Modellen und Sprachen zu identifizieren. Die Ergebnisse dieser Analyse wurden durch numerische Vergleichswerte ergänzt, die die qualitative Bewertung vervollständigten.

Übersetzungstools

ChatGPT-5 von OpenAI überzeugt durch Sprachgewandtheit, Kontextverständnis und Anpassungsfähigkeit in der Übersetzung. Besonders bei ressourcenstarken Sprachen liefert es flüssige und kohärente Ergebnisse, auch bei längeren Texten (Jiao et al. 2023; AlAfnan 2024; Wang et al. 2023). Durch Prompt Engineering kann es flexibel auf verschiedene Domänen reagieren, zeigt jedoch Schwächen bei geringen Ressourcen und komplexer Terminologie, was eine menschliche Nachbearbeitung erforderlich macht (Sahari et al. 2023; Siu 2023).

Gemini von Google zeichnet sich durch hohe Präzision bei idiomatischen, syntaktisch komplexen und ressourcenarmen Sprachen aus (Gupta et al. 2024; Nair/Gupta 2024; Jageer 2025). Seine Architektur legt den Fokus auf Genauigkeit und Fairness, wenngleich Sicherheitsfilter gelegentlich zu Auslassungen führen (Tekgurler 2025).

Perplexity ist einfacher, dafür aber schnell und prägnant. Es eignet sich für grundlegende Übersetzungs- und Lernaufgaben, erreicht jedoch nicht die Kontexttiefe von ChatGPT-5 oder Gemini (Sabrina et al. 2024).

4 Ergebnisse

Das folgende Kapitel präsentiert die zentralen Ergebnisse der Untersuchung. Dabei werden die Leistungen der drei Sprachmodelle ChatGPT-5, Gemini und Perplexity in Bezug auf die Übersetzungsqualität und den Erhalt metaphorischer Bedeutung systematisch ausgewertet und verglichen.

Insgesamt zeigten alle drei Systeme eine hohe Übersetzungsleistung (siehe Tab. 1). Perplexity lieferte dabei die flüssigsten, authentischsten und präzisesten Ergebnisse, insbesondere im slowakischen Testfall. ChatGPT-5 und Gemini erzielten ähnlich gute semantische und stilistische Resultate, diese wirkten jedoch etwas weniger natürlich. Die Durchschnittswerte über alle Metriken lagen hoch ($M = 4,3\text{--}4,9 / 5$), was auf eine generell starke Leistung aller Modelle hinweist.

Tab. 1: Deskriptive und inferenzstatistische Übersicht der Übersetzungsqualitätsmetriken

Metrik	Sprache	ChatGPT-5 M (SD)	Gemini M (SD)	Perplexity M (SD)	Haupteffekt: Tool (F, p, η^2_p)	Haupteffekt: Sprache (F, p, η^2_p)	Interaktion (F, p, η^2_p)
Flüssig- keit	EN / SK	4,36 (.55) / 4,52 (.48)	4,32 (.51) / 4,34 (.52)	4,42 (.49) / 4,92 (.15)	5,42; ,006; ,10	1,21; ,28; ,02	2,03; ,14; ,04
Ver- ständ- lichkeit	EN / SK	4,40 (.54) / 4,46 (.50)	4,40 (.51) / 4,52 (.47)	4,84 (.36) / 4,84 (.33)	2,14; ,12; ,04	4,89; ,032; ,09	1,15; ,32; ,02
Genauig- keit	EN / SK	4,30 (.53) / 4,50 (.49)	4,46 (.56) / 4,28 (.57)	4,38 (.54) / 4,88 (.36)	1,79; ,18; ,03	1,67; ,20; ,03	0,92; ,41; ,02
Authen- tizität	EN / SK	4,38 (.49) / 4,52 (.48)	4,34 (.53) / 4,46 (.51)	4,30 (.56) / 4,88 (.36)	6,17; ,003; ,11	2,64; ,11; ,05	3,12; ,049; ,06
Gesamt- index	EN / SK	4,36 (.53) / 4,50 (.48)	4,38 (.52) / 4,40 (.50)	4,49 (.46) / 4,88 (.31)	8,23; <,001; ,15	5,91; ,019; ,11	4,02; ,021; ,08

Signifikante Unterschiede zwischen den Tools zeigten sich bei Flüssigkeit ($F(2, 98) = 5,42$, $p = 0,006$, $\eta^2_p = 0,10$) und Authentizität ($F(2, 98) = 6,17$, $p = 0,003$, $\eta^2_p = 0,11$), wobei Perplexity konstant vorn lag. Ein Spracheneffekt wurde für Verständlichkeit ($F(1, 49) = 4,89$, $p = 0,032$) und den Gesamtindex ($F(1, 49) = 5,91$, $p = 0,019$) festgestellt. Dies legt etwas höhere Bewertungen für slowakische Übersetzungen nahe. Zudem ergaben sich Interaktionen zwischen Tool und Sprache bei Authentizität und Gesamtqualität, die den Vorsprung von Perplexity im Slowakischen betonen. Post-hoc-Analysen bestätigten Perplexity > ChatGPT-5 ($p < 0,01$) sowie Perplexity > Gemini ($p < 0,05$). Zwischen ChatGPT-5 und Gemini wurden keine signifikanten Unter-

schiede festgestellt. Die Effektgrößen ($\eta^2_p \approx 0,10-0,18$) deuten auf moderate, aber konsistente Vorteile von Perplexity hin.

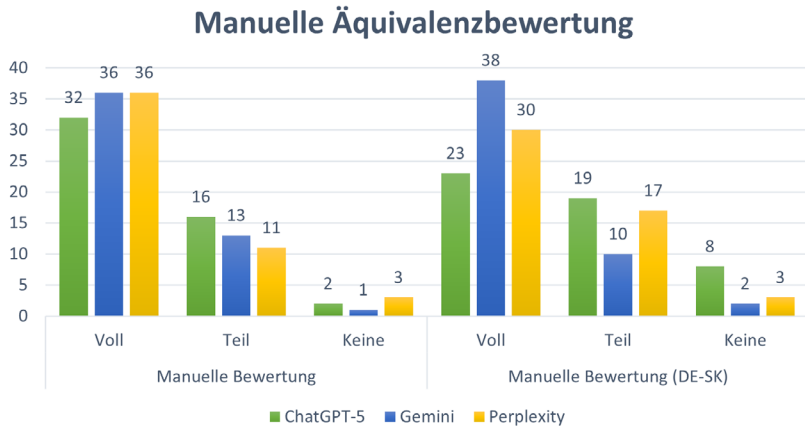


Abb. 1: Vergleich der manuellen Äquivalenzbewertung nach Übersetzungstool

Die manuelle Bewertung der Äquivalenz ergab unterschiedliche Leistungsmuster bei den drei großen Sprachmodellen (siehe Abb. 1). In der Sprachrichtung Deutsch–Slowakisch (DE–SK) verbesserte sich die Äquivalenzleistung insgesamt. Gemini lieferte die vollständigsten äquivalenten Übersetzungen (38), gefolgt von Perplexity (30) und ChatGPT-5 (23). Teilweise Äquivalenz trat häufig, aber zweitrangig auf, während Nicht-Äquivalenz nur sporadisch auftrat.

Insgesamt zeigen die Ergebnisse, dass Gemini durchweg den höchsten Anteil an vollständiger Äquivalenz erzielte, insbesondere bei den slowakischen Übersetzungen. Perplexity wiederum erzielte ein starkes Gleichgewicht zwischen vollständiger und teilweiser Äquivalenz. ChatGPT-5 lieferte mehr teilweise äquivalente Ergebnisse, was auf etwas wörtlichere oder kontextabhängigere Wiedergaben hindeutet.

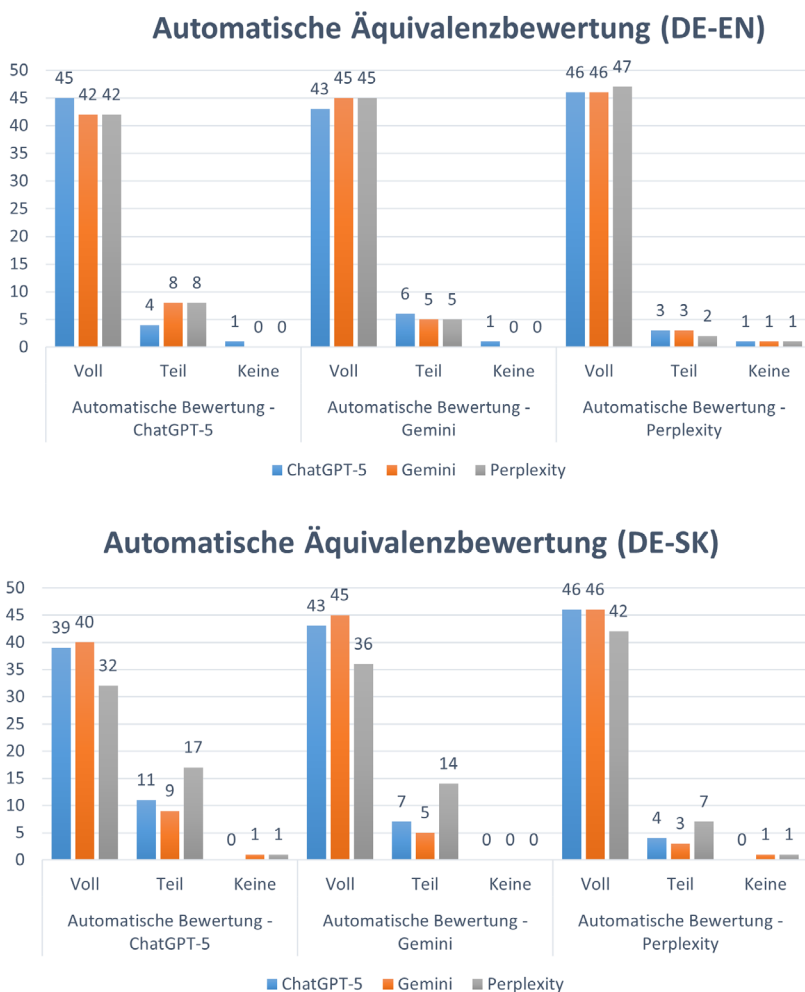


Abb. 2: Vergleich der automatischen Äquivalenzbewertung der Übersetzungstools ChatGPT-5, Gemini und Perplexity für die Sprachrichtungen Deutsch–Englisch und Deutsch–Slowakisch

Die automatischen Äquivalenzbewertungen (Abb. 2) zeigten fast durchgehend vollständige Übereinstimmung zwischen den drei LLMs in beiden Übersetzungsrichtungen ($\chi^2(4) = 1,06$, $p = 0,90$, $V = 0,07$ für DE–EN; $\chi^2(4) = 1,34$, $p = 0,85$, $V = 0,08$ für DE–SK). Bei Deutsch–Englisch stufen alle Modelle

über 40 von 50 Sätzen als vollständig äquivalent ein (ChatGPT-5 = 45, Gemini = 45, Perplexity = 47). Die hohen Cohen- κ -Werte (0,86–0,92) belegen eine nahezu perfekte Übereinstimmung. Ein ähnliches Muster ergab sich für Deutsch–Slowakisch, mit über drei Vierteln volläquivalenter Übersetzungen (Gemini = 45, Perplexity = 46; κ = 0,81–0,88).

Die manuellen Bewertungen zeigten dagegen mehr Variation ($\chi^2(4) = 2,18$, $p = 0,70$ für DE–EN; $\chi^2(4) = 6,92$, $p = 0,14$ für DE–SK), was auf differenziertere Beurteilungskriterien hinweist. Insgesamt neigen automatische Systeme dazu, vollständige Äquivalenz zu überschätzen, da sie subtile metaphorische oder kontextuelle Unterschiede weniger sensibel erfassen.

Interpretation ausgewählter Übersetzungen durch LLMs

Die Übersetzung des Ausdrucks „*argumentatives Futter für Zinssenkungen*“ (siehe Tab. 2) verdeutlicht die unterschiedlichen Strategien der drei getesteten Sprachmodelle im Umgang mit Metaphern. Die deutsche Formulierung kombiniert ein bildliches Element („Futter“) mit einem abstrakten wirtschaftlichen Fachwort („Zinssenkungen“), was eine Übertragung erfordert, die sowohl die metaphorische als auch die semantische Ebene wahrt.

Tab. 2: Qualitativer Vergleich der metaphorischen Übersetzungsstrategien dreier LLMs

Übersetzungstool	Ein schwacher Arbeitsmarkt liefert hierfür aber <i>argumentatives Futter für Zinssenkungen</i> .	
	EN	SK
ChatGPT-5	A weak labor market, however, provides <i>argumentative fuel for rate cuts</i> .	Slabý trh práce však poskytuje <i>argumentačnú muníciu pre zníženie sadzieb</i> .
Gemini	A weak labor market provides <i>argumentative fodder for interest rate cuts</i> .	Slabý trh práce však na to poskytuje <i>argumentačnú muníciu/podklad pre znížovanie úrokových sadzieb</i> .

Übersetzungs- tool	Ein schwacher Arbeitsmarkt liefert hierfür aber <i>argumentatives Futter für Zinssenkungen</i> .	
	EN	SK
Perplexity	A weak labor market provides <i>argumentative fodder for rate cuts</i> .	Slabý pracovný trh však poskytuje <i>argumenty pre znižovanie úrokov</i> .

ChatGPT-5 übersetzt den Ausdruck als „argumentative fuel for rate cuts“ und erhält damit die metaphorische Struktur in leicht abgeschwächter Form. Der Begriff *fuel* transportiert die Idee von „Antrieb“ oder „Brennstoff“ und wirkt idiomatisch sowie kontextuell passend. Somit gelingt eine sinngemäße und natürlich klingende Übertragung.

Gemini wählt die Formulierung „argumentative fodder for interest rate cuts“, was das Bild des „Futters“ direkter übernimmt, im ökonomischen Diskurs jedoch etwas unidiomatisch wirkt. Dennoch bleibt die Bildhaftigkeit gewahrt. Die slowakische Übersetzung („argumentačnú muníciu/podklad pre znižovanie úrokových sadzieb“) überträgt die Metapher treffend durch die Verwendung von „Munition“ als semantisch verwandtes Bild.

Perplexity entscheidet sich ebenfalls für „argumentative fodder for rate cuts“, verzichtet jedoch auf *interest* und erzeugt dadurch eine flüssigere, aber weniger bildhafte Variante. In der slowakischen Fassung („argumenty pre znižovanie úrokov“) wird die Metapher vollständig in eine sachlich-rationale Form überführt, was zwar die Bedeutung erhält, aber den figurativen Charakter aufgibt.

5 Diskussion und Zusammenfassung

In dieser Studie wurden metaphorische und idiomatische Ausdrücke im wirtschaftlichen Diskurs mithilfe von drei großen Sprachmodellen (LLMs) – ChatGPT-5, Gemini und Perplexity – übersetzt. Ein Vergleich der Ergebnisse für die Sprachpaare Deutsch-Englisch (mit umfangreichen Ressourcen) und

Deutsch-Slowakisch (mit geringen Ressourcen) ergab deutliche Unterschiede in der Übersetzungsqualität und im Umgang mit bildlicher Bedeutung.

Insgesamt zeigen die Ergebnisse, dass ChatGPT-5 und Gemini das Modell Perplexity in Bezug auf die Erhaltung der metaphorischen Bedeutung, der Sprachgewandtheit und der kontextuellen Genauigkeit übertreffen. Während Perplexity häufig wörtliche oder zu stark vereinfachte Wiedergaben erzeugte, was auf eine begrenzte Exposition gegenüber komplexen bildlichen Daten zurückzuführen ist, lieferten beide proprietären Modelle kohärente und semantisch präzise Übersetzungen. Diese Ergebnisse bestätigen Hypothese₁ und zeigen, dass proprietäre LLMs eine höhere Übersetzungsqualität liefern als frei zugängliche Systeme.

Beim Vergleich der Sprachressourcen erzielten Übersetzungen mit dem ressourcenreichen Sprachpaar (Deutsch–Englisch) eine bessere metaphorische Äquivalenz und stilistische Konsistenz als Übersetzungen mit dem ressourcenarmen Sprachpaar (Deutsch–Slowakisch). Dies stützt Hypothese₂. Die Wechselwirkung zwischen Modelltyp und Ressourcenebene (H₃) zeigte außerdem, dass Gemini unter ressourcenarmen Bedingungen relativ besser abschneidet, während ChatGPT-5 über beide Sprachpaare hinweg eine konsistente Qualität beibehält.

Diese Ergebnisse unterstreichen die anhaltende Datenabhängigkeit von LLM-basierten Übersetzungen, insbesondere bei idiomatischen und kulturell verankerten Inhalten. Obwohl LLMs die Übersetzung bildlicher Sprache verbessert haben, wird ihre Leistungsfähigkeit nach wie vor durch Datenvielfalt, kulturelle Nuancen und die Knappheit idiomatischer Ausdrücke in Trainingskorpora eingeschränkt.

Zukünftige Forschungsarbeiten sollten einen größeren und vielfältigeren Satz von Metaphern aus verschiedenen Genres und Bereichen umfassen. So lässt sich die statistische Validität stärken und das Verständnis von Metaphern umfassender testen. Darüber hinaus könnten gezielte Feinabstimmungen, erweiterte idiomatische und ressourcenarme Datensätze sowie die Entwicklung spezieller Bewertungsrahmen für die Übersetzung bildlicher Sprache zu weiteren Verbesserungen führen. Auch die Integration von multimodalem und kontextbezogenem Lernen könnte die Fähigkeit von LLMs verbessern, Metaphern natürlicher zu interpretieren.

Zusammenfassend lässt sich sagen, dass LLMs einen wichtigen Schritt in Richtung einer kulturell und semantisch bewussten Übersetzung darstellen. Eine kontinuierliche Anpassung und vielfältigere, auf Metaphern fokussierte Daten sind jedoch entscheidend, um eine wirklich menschenähnliche Leistung zu erreichen.

Literatur

- ABDELGAFFAR, M.A. / EL MOGY, A. / SHARAF, N.A. (2022): „Adapting Large Multilingual Machine Translation Models to Unseen Low-Resource Languages via Vocabulary Substitution and Neuron Selection.“ *Proceedings of the 15th Biennial Conference of the Association for Machine Translation in the Americas (Vol. 1: Research Track)*. Eds. D. Kuh / F. Guzmán. Orlando. 287–297.
- ALAFNAN, MOHAMMAD AWAD (2024): „Large Language Models as Computational Linguistics Tools: A Comparative Analysis of ChatGPT and Google Machine Translations.“ *Journal of Artificial Intelligence and Technology* 5: 20–32. DOI: 10.37965/jait.2024.0549.
- CASSANO, F. / GOUWAR, J. / LUCCHETTI, F. / SCHLESINGER, C. / FREEMAN, A. / ANDERSON, C.J. / FELDMAN, M.Q., / GREENBERG, M. / JANGDA, A. / GUHA, A. (2023): „Knowledge Transfer from High-Resource to Low-Resource Programming Languages for Code LLMs.“ *Proceedings of the ACM on Programming Languages* 8: 677–708. DOI: 10.1145/3689735.
- CASTALDO, ALESSANDRO / MONTI, JOHANNA (2024): „Prompting Large Language Models for Idiomatic Translation.“ *Proceedings of the 1st Workshop on Creative-Text Translation and Technology*. 37–44.
- DONTHI, S. / SPENCER, M. / PATEL, O. / DOH, J.Y. / RODAN, E. / ZHU, K. / O'BRIEN, S. (2024): „Improving LLM Abilities in Idiomatic Translation.“ *Proceedings of the First Workshop on Language Models for Low-Resource Languages*. 175–181. DOI: 10.48550/arxiv.2407.03518.
- GAIN, B. / BANDYOPADHYAY, D. / EKBAL, A. (2025): „Bridging the Linguistic Divide: A Survey on Leveraging Large Language Models for Machine Translation.“ *arXiv preprint*, arXiv:2504.01919. DOI: 10.48550/arxiv.2504.01919.
- GAO, H. / ZHANG, J. / ZHANG, P. / YANG, C. (2025): „Consistency Rating of Semantic Transparency: an Evaluation Method for Metaphor Competence in Idiom Under-

- standing Tasks.“ *Proceedings of the 31st International Conference on Computational Linguistics*, Abu Dhabi, UAE. 10460–10471.
- GOOGLE DEEPMIND (2025): *Gemini Advanced* [Large language model]. Google. <https://gemini.google.com/> (13.11.2025).
- GÓMEZ, S. / DOMINGO, M. / CASACUBERTA, F. (2024): „Interactive Machine Translation with Large Language Models in Low Resources Languages.“ *IberSPEECH* 2024.14: 66–70. DOI: 10.21437/iberspeech.2024-14.
- GUO, P. / REN, Y. / HU, Y. / LI, Y. / ZHANG, J. / ZHANG, X. / HUANG, H. (2024): „Teaching Large Language Models to Translate on Low-resource Languages with Textbook Prompting.“ *LREC-COLING* 2024: 15685–15697.
- GUPTA, A. / RASTOGI, A. / MALHOTRA, H. / RANGARAJAN, K. (2024): „Comparative Evaluation of Large Language Models for Translating Radiology Reports into Hindi.“ *The Indian Journal of Radiology & Imaging* 35: 88–96. DOI: 10.1055/s-0044-1789618.
- HE, Z. / LIANG, T. / JIAO, W. / ZHANG, Z. / YANG, Y. / WANG, R. / TU, Z. / SHI, S. / WANG, X. (2023): „Exploring Human-Like Translation Strategy with Large Language Models.“ *Transactions of the Association for Computational Linguistics* 12: 229–246. DOI: 10.1162/tacl_a_00642.
- ICHEN, N. / STAMENKOVIĆ, D. / HOLYOAK, K. (2023): „Large Language Model Displays Emergent Ability to Interpret Novel Literary Metaphors.“ *Metaphor and Symbol* 39: 296–309. DOI: 10.1080/10926488.2024.2380348.
- JAGEER, SYEDKHALEEL (2025): „A Comparative Analysis of Large Language Models for English-to-Tamil Machine Translation with Performance Evaluation.“ *Panamerican Mathematical Journal* 35.2s: 728–742. DOI: 10.52783/pmj.v35.i3s.3232.
- JIAO, W. / WANG, W. / HUANG, J. / WANG, X. / SHI, S. / TU, Z. (2023): „Is ChatGPT A Good Translator? Yes With GPT-4 As The Engine.“ *arXiv preprint*, arXiv:2301.08745 (21.11.2025).
- KALAŠ, FILIP (2025a): „Evaluation of Artificial Intelligence Translation.“ *Dynamics of Translation Studies – Digitization, Training, and Evaluation*. Berlin: Frank & Timme. 13–25.
- KALAŠ, FILIP (2025b): „Evaluation of German-Slovak AI Translation of Stock Market News.“ *Linguistische Treffen in Wrocław* 27.1: 117–129.
- KALAŠ, FILIP / LIPTÁK, PAVOL (2025): „AI and the Translation of Idioms – Challenges, Success, and a Corpus Perspective.“ *Jazykovedný časopis* 76.1: 258–267.

- KHOSHATAB, P. / NAMAZIFARD, D. / MASOUDI, M. / AKHGARY, A. / SANI, S. M. / YAGHOOBZADEH, Y. (2024): „Comparative Study of Multilingual Idioms and Similes in Large Language Models.“ *arXiv preprint*, arXiv:2410.16461. DOI: 10.48550/arxiv.2410.16461.
- LI, J. / ZHOU, H. / HUANG, S. / CHEN, S. / CHEN, J. (2023): „Eliciting the Translation Ability of Large Language Models via Multilingual Finetuning with Translation Instructions.“ *Transactions of the Association for Computational Linguistics* 12: 576–592. DOI: 10.1162/tacl_a_00655.
- LI, S. / CHEN, J. / YUAN, S. et al. (2023): „Translate Meanings, Not Just Words: IdiomKB's Role in Optimizing Idiomatic Translation with Language Models.“ *arXiv preprint*, arXiv:2308.13961. DOI: 10.48550/arxiv.2308.13961.
- LIANG, X. / KHAW, Y. / LIEW, S. / TAN, T. / QIN, D. (2025): „Toward Low-Resource Languages Machine Translation: A Language-Specific Fine-Tuning With LoRA for Specialized Large Language Models.“ *IEEE Access* 13: 46616–46626. DOI: 10.1109/access.2025.3549795.
- MOSLEM, Y. / HAQUE, R. / WAY, A. (2023): „Adaptive Machine Translation with Large Language Models.“ *arXiv preprint*, arXiv:2301.13294. DOI: 10.48550/arxiv.2301.13294.
- NAIR, A. / GUPTA, D. (2024): „Evaluating Performance and Accuracy of Large Language Models in Translating Code-Mixed Hindi to English: A Comparative Study.“ *2024 IEEE 21st India Council International Conference (INDICON)*. 1–6. DOI: 10.1109/indicon63790.2024.10958340.
- OPENAI (2025): *ChatGPT-5* [Large language model]. OpenAI. <https://chat.openai.com/> (13.11.2025).
- PERPLEXITY AI (2025): *Perplexity AI* [Artificial intelligence search and language model]. <https://www.perplexity.ai/> (13.11.2025).
- QIAN, MING (2024): „Automating Idiom Translation with Cross-Lingual Natural Language Generation Grounded In Semantic Analyses Using Large Language Models.“ *Proceedings of the 16th Conference of the Association for Machine Translation in the Americas (Vol. 2: Presentations)*. Chicago, IL: Association for Machine Translation in the Americas. 95–115.
- SABRINA, S. / SHAFI, C. / AINUN, A. / ASKIA, N. / MAISARAH, P. / SUSILAWATI, P. (2024): „Effectiveness of ChatGPT vs. Perplexity AI in English Language Lear-

- ning at the University of Serambi Mekkah.“ *IC-ITECHS*. DOI: 10.32664/ic-itechs.v5i1.1611.
- SAHARI, Y. / AL-KADI, A. / ALI, J. (2023): „A Cross Sectional Study of ChatGPT in Translation: Magnitude of Use, Attitudes, and Uncertainties.“ *Journal of Psycholinguistic Research* 52: 2937–2954. DOI: 10.1007/s10936-023-10031-y.
- SIU, SAI CHEONG (2023): „ChatGPT and GPT-4 for Professional Translators: Exploring the Potential of Large Language Models in Translation.“ *SSRN Electronic Journal*. DOI: 10.2139/ssrn.4448091.
- SONG, Y. / LI, L. / LOTHRTZ, C. / EZZINI, S. / SLEEM, L. / GENTILE, N. / STATE, R. / BISSYANDÉ, T. / KLEIN, J. (2025): „Is Small Language Model the Silver Bullet to Low-Resource Languages Machine Translation?“ *arXiv preprint*, arXiv:2503.24102. DOI:10.48550/arxiv.2503.24102.
- TAN, W. / ZHU, K. (2024): „NusaMT-7B: Machine Translation for Low-Resource Indonesian Languages with Large Language Models.“ *arXiv preprint*, arXiv:2410.07830. DOI: 10.48550/arxiv.2410.07830.
- TARS, M. / TÄTTAR, A. / FISHEL, M. (2022): „Cross-lingual Transfer from Large Multilingual Translation Models to Unseen Under-resourced Languages.“ *Baltic Journal of Modern Computing* 10. DOI: 10.22364/bjmc.2022.10.3.16.
- TEKGURLER, MERVE (2025): „LLMs for Translation: Historical, Low-Resourced Languages and Contemporary AI Models.“ *arXiv preprint*, arXiv:2503.11898. DOI: 10.48550/arxiv.2503.11898.
- TRAN, K. / O’SULLIVAN, B. / NGUYEN, H. (2024): „Irish-based Large Language Model with Extreme Low-Resource Settings in Machine Translation.“ *Proceedings of the Seventh Workshop on Technologies for Machine Translation of Low-Resource Languages (LoResMT 2024)*. 193–202. DOI: 10.18653/v1/2024.loresmt-1.20.
- WANG, L. / LYU, C. / JI, T. / ZHANG, Z. / YU, D. / SHI, S. / TU, Z. (2023): „Document-Level Machine Translation with Large Language Models.“ *arXiv preprint*, arXiv:2304.02210. DOI: 10.48550/arxiv.2304.02210.
- WANG, S. / ZHANG, G. / WU, H. / LOAKMAN, T. / HUANG, W. / LIN, CH. (2024): „MMTE: Corpus and Metrics for Evaluating Machine Translation Quality of Metaphorical Language.“ *arXiv preprint*, arXiv:2406.13698. DOI: 10.48550/arxiv.2406.13698.

- WANG, Y. / ZHANG, J. / SHI, T. / DENG, D. / TIAN, Y. / MATSUMOTO, T. (2024): „Recent Advances in Interactive Machine Translation With Large Language Models.“ *IEEE Access* 12: 179353–179382. DOI: 10.1109/access.2024.3487352.
- ZHU, W. / LIU, H. / DONG, Q. / XU, J. / HUANG, S. / KONG, L. / CHEN, J. / LI, L. (2023): „Multilingual Machine Translation with Large Language Models: Empirical Results and Analysis.“ *arXiv preprint*, arXiv:2304.04675. DOI: 10.48550/arxiv.2304.04675.
- ZHU, Y. / LAL, D. M. / DENYSIUK, S. / MITKOV, R. (2024): „From Neural Machine Translation to Large Language Models: Analysing Translation Quality of Chinese Idioms.“ *Proceedings of the International Conference on New Trends in Translation and Technology (INTTT 2024)*. 247–260. DOI: 10.26615/issn.2815-4711.2024_021.

Danksagung

Die Studie wurde im Rahmen des Projekts A-25-106/3020-18 „Künstliche Intelligenz: Herausforderungen für Linguistik und Marketing“ und des Projekts I-26-101-00 „Sprache als Verkaufsinstrument: Bildung von Vertrauen und Wirksamkeit in KI-gestützter Markenkommunikation“ durchgeführt.

Sprache als Verkaufsinstrument

Fahnenwörter im Marketing und ihre Auswirkungen auf den Verkauf im Zeitalter der KI

Abstract

The study examines the economic impact of linguistic flag words in online marketing communication. Using a dataset of 15,663 real sales across Slovak, Czech, and German platforms, the analysis combined AI-based text classification and manual validation to identify four linguistic categories: urgency, exclusivity, discount, and identity. Descriptive and inferential statistical methods, including Welch's t-test, were applied to evaluate price differences between products with and without flag words. The results confirm that linguistic elements significantly influence pricing strategies and consumer perception – identity and exclusivity increase perceived value, while discount and urgency support short-term sales dynamics. The study highlights the need for hybrid approaches combining Artificial Intelligence with human linguistic interpretation in digital marketing contexts (Kalaš/Lipták 2025).

Keywords AI and NLP in marketing, e-commerce, flag words, identity-based branding, marketing language

1 Einführung

Im Zeitalter der Digitalisierung und Künstlichen Intelligenz wird Sprache zu einem wichtigen wirtschaftlichen Instrument neben ihrer Funktion als Kommunikationsmittel. In einer Online-Umgebung, in der die Aufmerksamkeit

der Verbraucher nur für einen Bruchteil einer Sekunde gefangen wird, kommt der Formulierung des Produktnamens eine strategische Bedeutung zu. Die Verwendung sprachlicher Appelle, also von Wörtern, die den Eindruck von Dringlichkeit, Einzigartigkeit, Identität oder Vorteilhaftigkeit vermitteln, kann den wahrgenommenen Wert eines Produkts und die Kaufentscheidung der Verbraucher direkt beeinflussen.

Unser Ziel ist es, empirisch zu überprüfen, inwieweit sprachliche Appelle den Verkaufspreis, die Marktvariabilität und die wirtschaftliche Leistungsfähigkeit von Produkten im Online-Umfeld beeinflussen. Die Studie verbindet somit einen linguistischen und einen marketingorientierten Ansatz und zeigt das Potenzial Künstlicher Intelligenz bei der Identifizierung und Bewertung überzeugender Sprachstrategien auf.

2 Theoretische Ausgangspunkte

Die rasante Entwicklung der KI-Technologien verändert die digitale Kommunikation, insbesondere im Marketing, grundlegend (Chintalapati/Pandey 2021; Madanchian 2024). Sprache fungiert dabei nicht mehr nur als Informationsmittel, sondern als strategisches Instrument, das Wahrnehmung, Emotionen und Konsumententscheidungen prägt (Petty/Cacioppo 1986; Verma et al. 2021; Braca/Dondio 2023).

Im Online-Shopping spielt der Produktname eine zentrale Rolle. Erfolgreiche Bezeichnungen aktivieren psychologische Reize wie Knappheit, Dringlichkeit oder Exklusivität und fördern so die Kaufmotivation (Broeder/Wentink 2022; Zamfir 2024; Putriyandari et al. 2025). Studien bestätigen die starke Wirkung solcher Mangelimpulse: Wahrgenommene Verknappung erhöht die Attraktivität und Zahlungsbereitschaft (Hamilton et al. 2019; Ladeira et al. 2023; Roux et al. 2023), während übermäßige Dringlichkeit Misstrauen hervorrufen kann (Barton 2022; Rahmani 2023).

Sprache wird somit zu einem aktiven wirtschaftlichen Instrument, dessen Einfluss interdisziplinär untersucht wird (Kalaš/Lipták 2025; Lipták/Kalaš forthcoming). Laut Bhatia (2021) ist kommerzielle Kommunikation durch

das Zusammenspiel von Genre- und Sprachstrategien geprägt. Das Prinzip der Knappheit steigert den wahrgenommenen Wert eines Produkts und beschleunigt Kaufentscheidungen (Ladeira et al. 2023; Roux et al. 2023; Zamfir 2024). Zugleich hängt die Wirksamkeit solcher Appelle vom Vertrauen in den Verkäufer bzw. die Verkäuferin und dem wahrgenommenen Risiko ab (Lăzăroiu et al. 2020). Auch Identitätsaspekte beeinflussen Konsumententscheidungen. So stärkt die Betonung nationaler Herkunft („slowakisches Produkt“) emotionale Bindungen (Naylor et al. 2012), während eine klare visuelle und sprachliche Markenidentität positive Einstellungen und Zahlungsbereitschaft erhöht (Yu et al. 2024).

Durch die Weiterentwicklung von KI und NLP ergeben sich neue Möglichkeiten für eine automatisierte und personalisierte Marketingkommunikation (Chintalapati/Pandey 2021; Madanchian 2024). Sprachliche Preisschilder erweisen sich dabei als wirksames Mittel der Preisgestaltung (Lipták/Kalaš forthcoming), und NLP-Tools können produktoptimierte Beschreibungen mit „Flag Words“ erzeugen. Diese vermitteln Dringlichkeit oder Exklusivität und können so den Absatz steigern (Broeder/Wentink 2022; Jiang et al. 2023; Zou et al. 2023). Auf Grundlage der oben beschriebenen theoretischen Ausgangspunkte wurden für die vorliegende Studie folgende Hypothesen formuliert:

H₁: Sprachliche Appelle in Produktnamen haben einen statistisch signifikanten Einfluss auf deren wirtschaftliche Leistungsfähigkeit. Produkte mit sprachlichen Appellen erzielen im Vergleich zu Produkten ohne sprachliche Appelle höhere Verkaufszahlen und bessere Preise.

H₂: Die einzelnen Kategorien sprachlicher Appelle unterscheiden sich in ihrer wirtschaftlichen Wirkung. Appelle, die auf Identität und Exklusivität abzielen, tragen zu einem höheren Preis und einem höheren wahrgenommenen Wert des Produkts bei. Appelle, die auf Dringlichkeit und Rabatte abzielen, fördern einen schnelleren Verkauf und einen höheren Umsatz.

3 Methodologie

3.1 Beschreibung des Datensatzes und des Datenverarbeitungsprozesses

Für diese Studie wurde ein umfangreicher Datensatz von 15 663 tatsächlich verkauften Produkten aus einem slowakischen E-Shop, dem tschechischen Auktionsportal Aukro und der deutschen Version der Plattform eBay einer Analyse unterzogen. Die Daten wurden im Zeitraum von 2020 bis 2025 von einem einzigen Verkäufer erhoben. Auf der Grundlage dieser Daten wurde ein einheitliches, harmonisiertes Format für die anschließende analytische Verarbeitung des gesamten Datensatzes erstellt.

Tab. 1: Umfang des Datensatzes auf Basis der Plattformen

Plattform	Web	Anzahl der Artikel
Slowakische Verkaufsplattform	www.wdshop.eu	7 972
Tschechische Verkaufsplattform	www.aukro.cz	7 582
Deutsche Verkaufsplattform	www.ebay.de	109
Insgesamt		15 663

In dieser Studie wurde ein Datensatz verwendet, der bestimmte Produktkategorien – Sammler- und Investitionsgegenstände – umfasste. Diese Kategorie ist aus methodologischer Sicht einzigartig, da es sich nicht um eine homogene Produktgruppe mit einer einheitlichen Zielgruppe handelt, sondern um ein breites Spektrum von Artikeln mit unterschiedlich wahrgenommenem Wert und unterschiedlicher Verbrauchermotivation. Ein typisches Beispiel sind Anlagemünzen, bei denen die Materialzusammensetzung (Gold, Silber oder Kupfer) eine entscheidende Rolle spielt. Während Goldprodukte für Verbraucher mit höherer Kaufkraft bestimmt sind, richten sich Produkte aus Silber oder Kupfer an eine breitere Öffentlichkeit. Dieser Unterschied stellt jedoch kein methodisches Hindernis dar, sondern ermöglicht es, die Universalität des Spracheffekts über heterogene Verbrauchergruppen hinweg zu testen. Das Ziel der Studie ist es daher nicht, die Eigenschaften der Produkte oder ihre mate-

riellen Attribute zu untersuchen, sondern den Einfluss von Schlagworten in Produktnamen zu analysieren und zu quantifizieren, also zu untersuchen, wie sprachliche Elemente die Wahrnehmung des Werts und die Verkaufsfähigkeit von Produkten im Online-Umfeld beeinflussen.

Die Daten wurden in zwei Phasen verarbeitet: In der ersten Phase erfolgte eine automatisierte Textanalyse mit dem ChatGPT-5-Modell, in der zweiten Phase eine manuelle Überprüfung der Ergebnisse. Dieses Verfahren gewährleistet eine höhere Genauigkeit bei der Zuordnung von Sprachausdrücken zu Produkten und überprüft die Zuverlässigkeit des verwendeten Algorithmus. Der detaillierte Prozess der Datenvalidierung und seine Ergebnisse sind in Kapitel 4.1 „Datenvalidierung“ aufgeführt.

3.2 Statistische Methoden und Definition der Grundvariablen

Die wichtigste analysierte Variable in dieser Studie war der Stückpreis des Produkts. Er stellt einen Indikator für den wirtschaftlichen Wert dar und spiegelt gleichzeitig die Wertwahrnehmung der Verbraucher wider. Die Berechnungen wurden in Microsoft Excel (Version 2024) mithilfe des Moduls „Data Analysis Toolpak“ und der integrierten statistischen Funktionen AVERAGE, MEDIAN, MAX, MIN, QUARTILE und T.TEST durchgeführt. Die Ergebnisse dienten als Grundlage für die Interpretation der finanziellen Unterschiede zwischen den Kategorien sowie für die anschließende Diskussion über den Einfluss der Sprache auf das Verbraucherverhalten.

4 Ergebnisse und Diskussion

In diesem Kapitel wird auf die Verknüpfung von quantitativer Analyse und Interpretation der Ergebnisse hingewiesen, um die wirtschaftlichen und linguistischen Auswirkungen von Sprachfahnenwörtern im Online-Umfeld zu ermitteln. Die Ergebnisse werden in drei logischen Einheiten präsentiert:

Überprüfung der Genauigkeit der Datenverarbeitung mithilfe Künstlicher Intelligenz, Analyse der Finanzindikatoren und statistischer Vergleich der Kategorien.

4.1 Datenvalidierung

In dieser Phase wurde die Zuverlässigkeit der automatisierten Klassifizierung von Schlüsselwörtern überprüft, d. h., inwieweit KI in der Lage ist, Fahrenwörter in Marketingtexten ohne kontextuelle Interpretation selbstständig zu identifizieren. Die Validierung ergänzt somit unsere Studie und stellt einen Kontrollschritt dar, mit dem überprüft werden sollte, ob und inwieweit KI die manuelle linguistische Analyse in realen Daten ersetzen kann.

Die Ergebnisse des Vergleichs zwischen automatisierter und manueller Verarbeitung sind in der folgenden Tabelle aufgeführt. Sie fasst die Anzahl der richtigen und falschen Zuordnungen für die einzelnen Fahrenwort-Kategorien sowie den berechneten Zuverlässigkeitsgrad des ChatGPT-5-Modells zusammen.

Tab. 2: Vergleich der Sprachkategorien nach Art ihrer Auswahl

Kategorie	Anzahl der Vorkommen und deren Zuordnung (AI)			Manuelle Validierung	Zuverlässigkeit ChatGPT-5
	Σ	X	✓		
Urgenz	78	47	31	6060	0,51 %
Exklusivität	334	61	273	3027	9,02 %
Rabatt	536	2	534	1200	44,50 %
Identität	1060	444	616	624	98,72 %

Die Ergebnisse bestätigen, dass KI-Modelle, die mit allgemeinen Sprachkorpora arbeiten, spezifische Marketingbedeutungen nicht zuverlässig identifizieren können – unabhängig von grammatikalischer Korrektheit oder lexikalischer Ähnlichkeit. Für eine genaue Analyse ist es daher wichtig, die automatisierte Verarbeitung mit der menschlichen Interpretation zu kombinieren – ein so-

nannter hybrider Ansatz, der die Effizienz der KI mit dem Kontextverständnis eines Marketingexperten verbindet.

4.2 Finanzkennzahlen

Nach Abschluss des Validierungsprozesses wurde eine quantitative Analyse durchgeführt, um zu bewerten, ob das Vorhandensein von Sprachfahnenwörtern in Produktnamen deren Finanzkennzahlen, insbesondere den Stückpreis und den Verkaufswert, beeinflusst.

Tab. 3: Grundlegende Finanzkennzahlen nach Kategorien

Kategorie	Σ	Finanzkennzahlen (v €)			
		$\bar{x}$	$x_{\min}$	$x_{\max}$	$\tilde{x}$
Urgenz	6060	5,43	0,04	224,08	3,40
Exklusivität	3027	11,19	0,04	705,00	4,90
Rabatt	1200	5,38	0,16	145,00	4,50
Identität	624	21,33	0,04	705,00	7,35
–	6501	6,72	0,01	405,00	4,50

In diesem Fall ist der Vergleich zwischen Durchschnitt und Median wichtig, da er zeigt, inwieweit die Preisverteilung in einer bestimmten Kategorie durch Extremwerte (z. B. sehr teure Produkte) beeinflusst wird. Während der Durchschnitt alle Werte, einschließlich dieser Extreme, berücksichtigt, stellt der Median den typischen Preis dar, den der Verbraucher am häufigsten bezahlt.

In Tabelle 3 sieht man deutliche Unterschiede zwischen Durchschnitt und Median in allen fünf Kategorien. Der größte Unterschied ist in der Kategorie „Identität“ zu sehen: Der Durchschnitt liegt hier bei 21,33 €, der Median jedoch nur bei 7,35 €. Dies deutet darauf hin, dass der am häufigsten vorkommende Wert unter den Produkten, die die Identität des Produkts hervorheben, höher war als der Durchschnittswert der Produkte, denen keine Kategorie zugeordnet wurde.

Aus finanzieller Sicht sind auch Werte unterhalb der Referenzwerte wichtig. Diese betreffen Produkte ohne zugewiesene Kategorie, die für die Arbeit des Verkäufers mit weniger profitablen Produkten, die sich beispielsweise durch geringe Verbraucherinteressen auszeichnen, von Bedeutung sind. Trotz der niedrigeren Preise, zu denen Produkte aus den Kategorien „Urgenz“ und „Rabatt“ verkauft werden, kann man feststellen, dass diese Kategorien für den Handel wichtig sind, da sie die allgemeine finanzielle Stabilität des Händlers unterstützen.

Auf der Grundlage dieser Erkenntnisse wurde die Analyse um relative Indikatoren erweitert: die prozentuale Abweichung von der Referenzgruppe und den Preisindex. Diese Indikatoren spiegeln die relative Position der einzelnen Kategorien auf dem Markt besser wider.

Tab. 4: Erweiterte Finanzkennzahlen nach Kategorien

Kategorie	Σ	Finanzkennzahlen (in €)		Abweichung (%)	Index ($x_0 = 100$)
		$\bar{x}$	$\bar{x}_0$		
Urgenz	6060	5,43	6,72	– 19,2 %	≈ 81
Exklusivität	3027	11,19	6,72	+ 66,6 %	≈ 167
Rabatt	1200	5,38	6,72	– 19,9 %	≈ 80
Identität	624	21,33	6,72	+ 217,4 %	≈ 317

Um diese Problematik umfassender zu beleuchten, wurde die Abweichung von der Referenzgruppe analysiert, um zu untersuchen, welchen Einfluss Fahnennörter auf den durchschnittlichen Stückpreis der Produkte haben. Die Ergebnisse zeigen eindeutig, dass Fahnennörter das Interesse der Verbraucher:innen an bestimmten Produkten steigern können, was sich auch in einem höheren Preisniveau dieser Artikel widerspiegelt.

In diesem Zusammenhang sind Kategorien wichtig, die den Wert des Produkts in den Vordergrund stellen (zum Beispiel „exklusiv“, „limitiert“ oder „nationale Kennzeichnung“). Diese sprachlichen Elemente bringen dem Verkäufer einen deutlich höheren Wert: In der Kategorie „Identität“ liegt die Steigerung bei über 217 %, in der Kategorie „Exklusivität“ bei etwa 66 % im Vergleich zu Produkten ohne Fahnennwort.

Die Preisabweichung wurde anhand der folgenden Formel berechnet:

$$\text{Abweichung (\%)} = \frac{\bar{x}_{\text{Kategorie}} - \bar{x}_0}{\bar{x}_0} \times 100$$

Ein weiterer wichtiger Indikator ist der Preisindex. Er gibt an, um wie viel Prozent der Durchschnittspreis der Produkte einer bestimmten Kategorie im Vergleich zur Referenzgruppe höher oder niedriger ist. Die Referenzgruppe wurde in diesem Fall als ($\bar{x}_0 = 100$) definiert und für die Berechnung wurde folgende Formel verwendet:

$$\text{Preisindex} = \frac{\bar{x}_{\text{Kategorie}}}{\bar{x}_0} \times 100$$

Aus der Tabelle geht hervor, dass der Indexwert 317 in der Kategorie „Identität“ bedeutet, dass Produkte mit sprachlichem Appell im Durchschnitt mehr als dreimal so teuer sind wie Produkte ohne sprachlichen Appell. Die Kategorien „Urgenz“ und „Rabatt“ erzielen mit einem Index von etwa 80 hingegen niedrigere Durchschnittspreise, was ihren rationalen und kurzfristigen Charakter widerspiegelt.

Aus Marketing-Sicht zeigt der Index die Position der Kategorie in der Preishierarchie des Marktes an: Je höher der Index, desto stärker trägt der sprachliche Appell zur Premiumwahrnehmung des Produkts und zu seinem finanziellen Wert bei. Praktisch kann der Preisindex als Instrument zur Preissegmentierung genutzt werden. Er ermöglicht es Verkäufern, zu erkennen, welche Sprachkategorien den wahrgenommenen Wert eines Produkts steigern und welche einen schnellen Lagerumschlag fördern. Durch die Kombination von hochwertigen („Identität“, „Exklusivität“) und margenschwachen („Rabatt“, „Urgenz“) Kategorien kann das Produktportfolio optimiert und ein Gleichgewicht zwischen Rentabilität und Umsatz erreicht werden.

4.3 Visualisierung und statistischer Vergleich der Fahnenwörter-Kategorien

Um den Zusammenhang zwischen den Sprachkategorien und dem Preisniveau besser zu verstehen, wurde der Datensatz mithilfe eines Boxplots visualisiert. Dieser ermöglicht einen Vergleich der Preisverteilung der Produkte in den einzelnen Fahnenwort-Kategorien (Abb. 1).

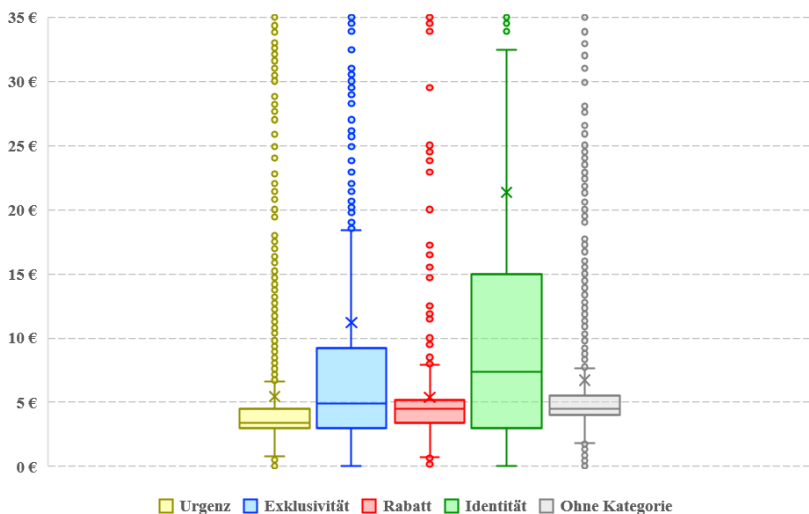


Abb. 1: Aufteilung der Produkte nach Fahnenwort-Kategorien

Bei Interpretation von Abb. 1 hat man sich auf die Interquartilsbreite (IQR) konzentriert. Diese stellt die Differenz zwischen dem dritten (Q3) und dem ersten Quartil (Q1) dar und gibt den Bereich an, in dem sich 50 % aller Beobachtungen befinden. Die Kategorie „Identität“ weist die größte IQR auf. Dies deutet auf eine erhebliche Preisvariabilität hin, da sie ein breites Spektrum an Produkten umfasst – von Standardartikeln bis hin zu Premium- und limitierten Produkten. Dieses Phänomen signalisiert einen großen Abstand zwischen dem unteren und dem oberen Quartil. Dieser kann auf mehrere

Faktoren zurückzuführen sein, z. B. auf limitierte Auflagen oder länderspezifische Produkte mit hoher Wertschöpfung.

Die Kategorie „Exklusivität“ hat zwar eine kleinere IQR als „Identität“, weist jedoch weiterhin eine Rechtsschiefe auf. Dies deutet auf eine Konzentration der meisten Preise im unteren Bereich und das Vorhandensein einiger deutlich teurerer Artikel hin. Somit weisen beide Kategorien einen langen „Schwanz“ der Verteilung in Richtung höherer Werte auf, was in der statistischen Terminologie als positive Schiefe (*rechtsschiefe Verteilung*) bezeichnet wird.

Aus statistischer Sicht ist es wichtig, dass die Unterschiede zwischen den einzelnen Quartilen größer sind als die Unterschiede zwischen den Medianwerten. Das bedeutet, dass Kategorien mit einem höheren IQR („Identität“, „Exklusivität“) eine größere interne Streuung aufweisen, während die Kategorien „Rabatt“ und „Urgenz“ eine höhere Preisstabilität zeigen. Eine solche Verteilung stützt die Hypothese, dass der sprachliche Appell nicht nur mit dem Preisniveau, sondern auch mit der Marktvariabilität zusammenhängt: Emotionale und wertbezogene Appelle führen zu einem preislich verstreuten Portfolio, während rationale Appelle zu einheitlicheren Preisen führen.

Angesichts der vorhandenen deutlichen Extreme musste überprüft werden, ob die beobachteten Unterschiede zwischen den Kategorien nicht nur auf eine variable Verteilung zurückzuführen sind, sondern auch statistisch signifikante Unterschiede in den Durchschnittswerten darstellen. Aus diesem Grund wurde anschließend eine Reihe von t-Tests unabhängiger Stichproben (Welch-t-Test) durchgeführt.

Tab. 5: Ergebnisse der t-Tests unabhängiger Stichproben (Student- und Welch-Test)

Kategorie	t-Test unabhängiger Stichproben			Auswirkung auf den Preis
	Student	Welch	Bedeutung	
Urgenz	$3,44 \times 10^{-8}$	$2,49 \times 10^{-8}$	$p < 0,001$	leicht positiv
Exklusivität	$1,40 \times 10^{-25}$	$1,08 \times 10^{-17}$	$p < 0,001$	stark konsistent
Rabatt	0,002268	$3,60 \times 10^{-6}$	$p < 0,01$	negativ
Identität	$1,99 \times 10^{-66}$	$1,59 \times 10^{-13}$	$p < 0,001$	sehr stark, variabel

Da alle p-Werte deutlich unter dem Signifikanzniveau $\alpha = 0,05$ liegen, ist die Wahrscheinlichkeit, dass die festgestellten Unterschiede zufällig sind, vernachlässigbar. Die Unterschiede zwischen den Ergebnissen des Student-Tests und des Welch-Tests sind auf ungleiche Streuung und Stichprobengrößen zurückzuführen. Dies bestätigt die Heterogenität der Daten. Hinsichtlich der Wirkung sind die Kategorien „Identität“ und „Exklusivität“ am stärksten. Neben einer hohen statistischen Signifikanz weisen sie auch eine deutliche Preisverschiebung gegenüber der Referenzgruppe auf.

5 Zusammenfassung

Die Ergebnisse der Analyse bestätigen, dass sprachliche Fahnenwörter einen messbaren Faktor darstellen, der das Verbraucherverhalten beeinflusst. Auf dieser Grundlage eröffnet sich ein neues Anwendungsgebiet für Künstliche Intelligenz: Sie kann nicht nur als analytisches Werkzeug zur Textverarbeitung dienen, sondern auch als aktiver Gestalter von Sprachstrategien.

Im Fall von H_1 wurde bewiesen, dass das Vorhandensein eines sprachlichen Appells im Produktnamen dessen wirtschaftliche Leistungsfähigkeit statistisch signifikant erhöht: Produkte mit Appellen wiesen entweder höhere Durchschnittspreise oder eine bessere Verkaufbarkeit als die Kontrollgruppe auf.

Die Analyse von H_2 hat zudem ergeben, dass sich die einzelnen Arten von Appellen in ihrer Funktion im Verkaufsprozess unterscheiden: Appelle, die auf Identität und Exklusivität ausgerichtet sind, tragen zur Wertschöpfung und zur Premiumpositionierung des Produkts bei. Appelle, die auf Dringlichkeit und Rabatte abzielen, stimulieren dagegen einen schnelleren Umsatz und kurzfristige Verkäufe.

Künstliche-Intelligenz-Modelle könnten in Zukunft dazu verwendet werden, Produkten automatisch Fahnenwörter zuzuordnen und optimierte Artikelbezeichnungen zu generieren. Diese würden gezielt die Präferenzen des Verkäufers oder die Marketingstrategie der Plattform widerspiegeln. In der Praxis könnte der Händler beispielsweise den Zielappell (Urgenz, Exklusivität, Rabatt oder Identität) definieren, während die KI die Sprache des Produkts

entsprechend anpassen würde, um dessen Sichtbarkeit und Attraktivität im Online-Umfeld zu maximieren. Ein solches System würde eine dynamische sprachliche Steuerung des Vertriebs ermöglichen: Produktnamen könnten sich je nach Saison, Zielgruppe oder aktueller Leistung der Kategorie ändern. Im Gegensatz zum derzeitigen statischen Ansatz würde Sprache so zu einem flexiblen Marketinginstrument werden, das in Echtzeit auf Daten reagiert.

Die Kombination aus Datenanalyse, NLP-Modellen und Marketingstrategie ermöglicht somit ein neues Konzept des automatisierten Sprachmarketings. Dabei interpretiert und gestaltet Künstliche Intelligenz Sprache aktiv – im Einklang mit den Zielen des Verkäufers und den Erwartungen der Verbraucher.

Literatur

- BHATIA, VIJAY K. (2021): *Critical Genre Analysis: Investigating Interdiscursive Performance in Professional Practice*. London / New York: Routledge. DOI: 10.4324/9781315690315.
- BRACA, ALESSIA / DONDIO, PIERPAOLO (2023): „Developing Persuasive Systems for Marketing: The Interplay of Persuasion Techniques, Customer Traits and Persuasive Message Design.“ *Italian Journal of Marketing* 2023: 369–412. DOI: 10.1007/s43039-023-00077-0.
- CHINTALAPATI, S. / PANDEY, S. K. (2021): „Artificial Intelligence in Marketing: A Systematic Literature Review.“ *International Journal of Market Research* 64.1: 38–68. DOI: 10.1177/14707853211018428.
- HAMILTON, R. / THOMPSON, D. V. / BONE, S. A. et al. (2019): „The Effects of Scarcity on Consumer Decision Journeys.“ *Journal of the Academy of Marketing Science* 47.3: 532–550. DOI: 10.1007/s11747-018-0604-7.
- JIANG, Y. / PANG, P. C.-L. / WONG, D. / KAN, H. Y. (2023): „Natural Language Processing Adoption in Governments and Future Research Directions: A Systematic Review.“ *Applied Sciences* 13.22: 12346. DOI: 10.3390/app132212346.
- KALAŠ, FILIP / LIPTÁK, PAVOL (2025): „AI and the Translation of Idioms – Challenges, Success, and a Corpus Perspective.“ *Journal of Linguistics* 76.1: 258–267.
- LADEIRA, W. J. / LIM, W. M. / SANTINI, F. D. O. et al. (2023): „A Meta-Analysis on the Effects of Product Scarcity.“ *Psychology & Marketing* 40: 1267–1279. DOI:10.1002/mar.21816.

- LĂZĂROIU, GEORGE / NEGURIȚĂ, OANA / GRECU, I. et al. (2020): „Consumers’ Decision-Making Process on Social Commerce Platforms: Online Trust, Perceived Risk, and Purchase Intentions.“ *Frontiers in Psychology* 11: 890. DOI:10.3389/fpsyg.2020.00890.
- LIPTÁK, PAVOL / KALAŠ, FILIP (forthcoming): „Language as a Sales Tool: The Power of Multilingual Pricing Tags.“
- MADANCHIAN, MITRA (2024): „The Impact of Artificial Intelligence Marketing on E-Commerce Sales.“ *Systems* 12.10: 429. DOI: 10.3390/systems12100429.
- BARTON, B. / ZLATEVSKA, N. / OPPEWAL, H. (2022): „Scarcity Tactics in Marketing: A Meta-Analysis of Product Scarcity Effects on Consumer Purchase Intentions.“ *Journal of Retailing* 98.4: 741–758. DOI: 10.1016/j.jretai.2022.06.003.
- NAYLOR, R. W. / LAMBERTON, C. P. / NORTON, D. A. (2012): „Seeing Ourselves in Others: Reviewer Ambiguity, Egocentric Anchoring, and Persuasion.“ *Journal of Marketing Research* 50.5: 617–631. DOI: 10.1509/jmkr.48.3.617.
- PETTY, RICHARD E. / CACIOPPO, JOHN T. (1986): *Communication and Persuasion: Central and Peripheral Routes to Attitude Change*. New York: Springer.
- PUTRIYANDARI, R. / AMRAN, A. / HARDJAYANTI, N. B. (2025): „Scarcity Strategies and Price Promotions in Enhancing Consumer Purchase Intensity.“ *International Journal of Management Science and Information Technology (IJMSIT)* 5.2: 303–310.
- RAHMANI, V. (2023): „Persuasion Knowledge Framework: Toward a Comprehensive Model of Consumers’ Persuasion Knowledge.“ *AMS Review* 13: 12–33. DOI: 10.1007/s13162-023-00254-6.
- ROUX, C. / GOLDSMITH, K. / CANNON, C. (2023): „On the Role of Scarcity in Marketing: Identifying Research Opportunities across the 5 Ps.“ *Journal of the Academy of Marketing Science* 51: 1197–1202. DOI: 10.1007/s11747-023-00956-0.
- BROEDER, P. / WENTINK, E. (2022): „Limited-time scarcity and competitive arousal in E-commerce.“ *The International Review of Retail, Distribution and Consumer Research* 32.5: 549–567. DOI: 10.1080/09593969.2022.2098360.
- SUOMALA, JARI (2020): „The Consumer Contextual Decision-Making Model.“ *Frontiers in Psychology* 11: 570430. DOI: 10.3389/fpsyg.2020.570430.
- VERMA, S. / SHARMA, R. / DEB, S. / MAITRA, D. (2021): „Artificial Intelligence in Marketing: Systematic Review and Future Research Direction.“ *International Journal of Information Management Data Insights* 1.2: 100002. DOI: 10.1016/j.jjimei.2020.100002.

- YU, M. / ZAINAL ABIDIN, S. B. / SHAARI, N. B. (2024): „Effects of Brand Visual Identity on Consumer Attitude: A Systematic Literature Review (Preprint).“ DOI: 10.20944/preprints202405.1109.v1.
- ZAMFIR, M. D. (2024): „Scarcity Effect and Consumer Decision Biases: How Urgency Influences the Perceived Value of Products.“ *Journal of World Economy* 3.4: 27–34. <https://www.pioneerpublisher.com/jwe/article/view/1095> (13.11.2025).
- ZOU, Y. / ZHANG, X. / ZHOU, J. et al. (2023): „Automatic Product Copywriting for E-Commerce.“ *AI Magazine* 44.1: 41–53. DOI: 10.1002/aaai.12084.

Danksagung

Die Studie wurde durchgeführt im Rahmen des Projekts A-25-106/3020-18 „Künstliche Intelligenz: Herausforderungen für Linguistik und Marketing“ und des Projekts I-26-101-00 „Sprache als Verkaufsinstrument: Bildung von Vertrauen und Wirksamkeit in KI-gestützter Markenkommunikation“ sowie des Projekts VEGA 1/0392/23 „Veränderungen im Ansatz zur Erstellung von Vertriebsmanagementkonzepten für Unternehmen unter dem Einfluss der Auswirkungen sozialer und wirtschaftlicher Krisen, verursacht durch die globale Pandemie und erhöhte Sicherheitsrisiken“.

KI-Einsatz zur Übersetzung deutsch-italienischer Wirtschaftstexte

Chance oder Risiko?

Abstract

This study examines the use of AI-assisted systems, such as DeepL, for translating German–Italian economic texts. The analysis is based on experiences from the Master’s programme *Language for Communication in International Enterprise and Organisation* at the University of Modena and Reggio Emilia (2023/24–2024/25), within the *German Specialised Communication and Translation* course. Initially, students translated a selection of corporate report passages independently. Subsequently, they translated the same passages with DeepL. The analysis reveals discrepancies in terminology accuracy, idiomatic and metaphorical expression use, and pragmatic nuances, comparing student and machine translations. This assertion is supported by literature on specialised translation (Biel/Sosoni 2017: 351–361) and AI-assisted translation (O’Brien/Simard/Specia 2012; Yang et al. 2023). The findings indicate that while machine translation enhances efficiency and produces idiomatic sentence structures, human expertise remains essential for metaphors, stylistic register, and culturally specific terms. The study emphasises the importance of integrating human expertise with AI tools, providing insights relevant to translator training and the incorporation of machine assistance into professional practice.

Keywords Fachübersetzung, maschinelle Übersetzung, Wirtschaftssprache

1 Einleitung

Die zunehmende Entwicklung der Künstlichen Intelligenz (KI) und ihr Gebrauch in allen menschlichen Bereichen tragen auch zu einer Veränderung von Übersetzungsprozessen bei, insbesondere was Fachsprachen betrifft. Letztere sind durch terminologische Dichte, hohe Kontextgebundenheit und kulturelle Spezifika gekennzeichnet (Wiesmann 2018). KI-gestützte Systeme wie ChatGPT, DeepL und Claude, um nur einige zu nennen, eröffnen neue Möglichkeiten zur Effizienzsteigerung: Sie können große Textmengen analysieren, terminologische Konsistenzen prüfen und Vorschläge für idiomatische Formulierungen geben. Gleichzeitig enthalten sie aber substantielle Herausforderungen hinsichtlich semantischer Präzision, Zuverlässigkeit sowie Transparenz und Nachvollziehbarkeit maschinell erzeugter Inhalte. Besonders in der Fachübersetzung wirtschaftlicher Texte erfordert dies ein vertieftes Nachdenken sowohl auf theoretischer als auch auf praktischer Ebene (Zaykova/Shilnikova 2019). Dies fügt sich in die aktuelle Debatte über die Rolle des Menschen im Übersetzungsprozess ein, die sich maßgeblich mit dem Konzept *human-in-the-loop translation* (Witczak 2016: 33–40; Yang et al. 2023; Hiebl/Gromann 2023: 375–378) befasst. Demzufolge wird die Interaktion zwischen Übersetzer:in und KI-System als kollaborativer Prozess verstanden, der spezifische Kompetenzen erfordert – insbesondere im Bereich der *post-editing competence* (O’Brien/Simard/Specia 2012; TAUS Guidelines 2017).

Dieser Beitrag untersucht die Potenziale und Grenzen von KI-gestützten Systemen im Übersetzungsprozess wirtschaftlicher Fachtexte. Die Überlegungen basieren auf praktischen Erfahrungen, die im Rahmen des Masterkurses *German Specialised Communication and Translation* im Studiengang *Language for Communication in International Enterprise and Organisation* (LACOM)¹ an der Universität Modena und Reggio Emilia in den Jahren 2023/2024, 2024/2025 gesammelt wurden. In diesem didaktischen Kontext wurde mit einer begrenzten Auswahl an KI-Tools experimentiert, um deren

1 Vgl. das Programm des Kurses: <https://unimore.coursecatalogue.cineca.it/insegnamento/2024/25250/2019/9999/10579?coorte=2023&schemaid=20897>, letzter Zugriff: 30.10.2025.

Nutzen für die Übersetzung komplexer Fachtexte zu testen und kritisch darüber nachzudenken.

Zwei Forschungsfragen leiten die folgende Analyse:

1. Inwieweit können KI-gestützte Systeme menschlichen Übersetzer:innen bei der Erstellung wirtschaftssprachlicher Fachtexte sinnvoll helfen?
2. Wo liegen die Grenzen solcher Systeme mit Bezug auf inhaltliche Tiefe, kulturelle Sensibilität und kommunikative Angemessenheit?

Die Untersuchung beabsichtigt, das Verhältnis zwischen Mensch und Maschine im Übersetzungsprozess zu betrachten und über die Rolle von KI nicht nur als technisches Werkzeug, sondern eher als möglicher „Dialogpartner“ im kognitiven Prozess des Übersetzungsprozesses zu reflektieren (Kiraly 2014: 1–14; Risku et al. 2017: 3–10).

2 Die Übersetzung wirtschaftlicher Texte

Die Übersetzung wirtschaftlicher Texte ist ein besonders komplexer Teilbereich der Fachübersetzung und stellt auf mehreren Ebenen spezifische Anforderungen: sprachlich-terminologisch, funktional-pragmatisch und kulturell-diskursiv (Biel/Sosoni 2017: 351–355; Lutz 2017: 284–292). Wirtschaftstexte wie Geschäftsberichte, Finanzanalysen, Pressemitteilungen oder wirtschaftspolitische Dokumente erfüllen grundlegend eine informative Funktion (Simonnæs/Kristiansen 2018: 151–160; Schropp/Dippoldt 2018: 135–149; Roelcke 2014: 154–160). Zugleich sind sie auch performativ, weil sie das Vertrauen von *Stakeholdern*, die Wahrnehmung von Institutionen und die Entscheidungsprozesse auf Märkten beeinflussen (Scarpa 2020: 57). Deshalb spielen Übersetzer:innen von Wirtschaftstexten eine bedeutende Rolle, die über die reine sprachliche Vermittlung hinausgeht, und werden zu aktiven Gestaltern ökonomischer Kommunikation.

Obwohl die terminologische Dichte wirtschaftlicher Fachtexte zunächst den Eindruck systematischer Eindeutigkeit vermittelt, zeigt eine vertiefte Analyse, dass gerade diese Dichte ein erhebliches Übersetzungsproblem verursacht. Die fachterminologische Präzision setzt sich mit der kultur- und rechtssystembedingten Polysemie vieler Schlüsselbegriffe auseinander (Stolze 2009: 148, 172).

Ein Beispiel dafür ist der englische Terminus *equity*: Im britischen Rechnungswesen bedeutet er meist „Eigenkapital“, kann aber in juristischen Kontexten auch „Gerechtigkeit“ oder „Billigkeit“ heißen (vgl. DGRV 2006 : 11f.). Die maschinelle Übersetzung, die auf distributionelle Semantikmodelle und statistische Kookkurrenzen zurückgreift, verfügt jedoch nicht über die außersprachlichen Kompetenzen, um konzeptuelle Unterschiede zu erfassen. Diesbezüglich stellt sich eine zentrale Frage für KI-gestützte Systeme: Inwieweit können neuronale Netzwerke jene Form Expertenwissen rekonstruieren, die nicht in den Oberflächenstrukturen von Trainingskorpora vermittelt wird, sondern in den Logiken der Fachdomäne selbst begründet ist?

Neben der terminologischen Komplexität ist wirtschaftliche Fachkommunikation von einer pragmatischen Dimension charakterisiert. Demzufolge sind Übersetzungen nicht primär unter dem Aspekt formaler Äquivalenz zu bewerten, sondern vielmehr hinsichtlich ihrer Adäquatheit im Zielkontext (Stolze 2009: 44–47, 63, 113–118; Lüger 2018: 71–75). Ein Geschäftsbericht erfüllt z. B. nicht nur informative Funktionen, sondern trägt zugleich zur Imagebildung und zu den *Investor Relations* bei. Somit werden Übersetzungsstrategien gebraucht, die über rein semantische Transferprozesse hinausgehen: Metaphern bzw. institutionell konventionalisierte Ausdrucksweisen müssen so adaptiert werden (Stolze 2009: 140–143), dass sie im Zielsystem verständlich bleiben, ohne die persuasive Intention des Ausgangstextes zu verlieren. Tcaciuc/Mackevic (2017: 471–433) zeigen z. B. in ihrer Analyse von EU-Wirtschaftstexten, dass Übersetzer:innen hierfür systematisch kreative Kompensationsstrategien verwenden, um sowohl die semantische als auch die pragmatische Textintegrität zu gewährleisten. Zudem wird das Verständnis von system- und kulturspezifischen Divergenzen, wie unterschiedlichen Rechnungslegungsstandards, Rechtstraditionen oder branchenspezifischen Diskurskonventionen, erschwert. Diese bilden asymmetrische Wissensräume,

die bei jeder Übersetzungsentscheidung mitverhandelt werden müssen. Zaykova/Shilnikova (2019) betonen ferner, dass Übersetzen im wirtschaftlichen Bereich ein stetiger Prozess ist, in dem Übersetzer:innen zwischen normativen Systemen und kulturellen Rahmenbedingungen Inhalte vermitteln. Dies wird von einem Begriff wie *shareholder value* besonders deutlich illustriert: Im anglo-amerikanischen Diskurs bezeichnet das Konzept ein Leitprinzip der Unternehmensführung, das auf die Maximierung des Aktionärsvermögens abzielt. Im deutschsprachigen Kontext ist der Terminus direkt übernommen worden, hat jedoch eine spezifische diskursive Konnotation entwickelt, die oft kritisch mit einer kurzfristigen Gewinnorientierung assoziiert wird (vgl. Janisch 1993: 114). Übersetzer:innen und KI-gestützte Übersetzungssysteme müssen sich solcher Unterschiede bewusst sein, um Fehlinterpretationen zu vermeiden. Nicht zu unterschätzen sind schließlich die stilistischen Besonderheiten ökonomischer Texte.

Die Sprache der Wirtschaft ist oft geprägt von Nominalstil, komplexen Satzgefügen, Anglizismen und konzeptuellen Metaphern (vgl. Lemon/Fluck 2016: 69–80). Letztere, wie z. B. die Vorstellung des Marktes als „unsichtbare Hand“ oder des Unternehmens als „Organismus“, vermitteln nicht nur Informationen, sondern strukturieren ganze Diskurse. Ihre Übersetzung stellt eine doppelte Herausforderung dar: Einerseits soll die metaphorische Ausdruckskraft sichergestellt werden; andererseits soll die Metapher auch im Zielkontext verstanden werden (Scarpa 2020: 230–240; Tcaciuc/Mackevic 2017: 417–420).

Die Übersetzung wirtschaftlicher Texte eröffnet ein breites Forschungsfeld, das mit Fragen nach der translatorischen Entscheidungspraxis, der Rolle institutioneller Normen sowie der Wirkung auf das Zielpublikum verbunden ist. In diesem Zusammenhang stellt der Einsatz von KI diese Form des Denkens vor neue Herausforderungen: Was passiert mit der kognitiven Leistung des Übersetzens, wenn maschinelle Vorschläge bereits große Teile des Textes „vorformulieren“? Die Beobachtungen aus den Erfahrungen im oben genannten Masterkurs beweisen, dass KI-gestützte Systeme nicht als Ersatz menschlicher Übersetzer:innen zu begreifen sind, sondern vielmehr als Werkzeug im Kontext eines interaktiven Übersetzungsprozesses.

3 Methodologie

Die vorliegende Untersuchung basiert auf einem praxisorientierten Ansatz, der im Rahmen des Masterkurses *German Specialised Communication and Translation* in den akademischen Jahren 2023/2024 und 2024/2025 entwickelt wurde. Ziel war es, die praktische Anwendbarkeit verschiedener KI-Tools für die Übersetzung wirtschaftlicher Fachtexte zu testen und kritisch zu überprüfen. Der Ansatz orientierte sich an Prinzipien der *empirical classroom research* bzw. *action research* (Rico / González Pastor 2022: 177–197). Die Datenerhebung und Analyse erfolgten unmittelbar im didaktischen Kontext des Masterkurses und verfolgten sowohl Lern- als auch Forschungsziele. Der methodische Zugang stützte sich auf eine vergleichende und qualitative Analyse von Übersetzungen und vollzog sich in einem zweistufigen Verfahren: In einem ersten Schritt wurde eine Auswahl von Textpassagen aus deutschen Geschäftsberichten von den Studierenden ins Italienische übersetzt. Dabei wurde auf den Einsatz KI-gestützter Systeme verzichtet. Stattdessen basierte die Übersetzung auf der Sprachkompetenz der Studierenden sowie auf der Nutzung traditioneller Nachschlagewerke (Fuertes-Olivera/Nielsen 2012: 191–215), Plattformen für Terminologie (IATE) und der Software *Sketch Engine*, um Okkurrenzen bzw. Kookkurrenzen zu prüfen (vgl. Zanettin 2018). In einem zweiten Schritt wurden dieselben Textpassagen mithilfe des KI-Übersetzungssystems DeepL übersetzt. DeepL wurde aufgrund seiner führenden Position im Bereich der maschinellen Übersetzung für europäische Sprachen und seiner besonderen Eignung für Fachtexte ausgewählt (Gerhalter 2024: 1–6). Im Anschluss daran wurden menschliche und maschinelle Übersetzungen systematisch miteinander verglichen, um Unterschiede in Stil, Wortwahl und Übersetzungsstrategien zu identifizieren. Dieser Ansatz ermöglichte es, sowohl die spezifischen Leistungen der KI-Tools zu bewerten als auch den Mehrwert menschlicher Übersetzungskompetenz sichtbar zu machen.

Die Analyse erfolgte in iterativen Überarbeitungszyklen, in denen sowohl die manuellen als auch die KI-gestützten Übersetzungen analysiert, kommentiert und bei Bedarf revidiert wurden. Besonderes Augenmerk lag dabei auf Fehlertypen, die sowohl für menschliche Übersetzungen als auch für KI-Systeme typisch sind, wie semantische Verschiebungen, kontextunangemessene

Termini, kulturell unpassende Formulierungen oder der Verlust performativer Funktionen.

4 Korpusmaterial

Das Korpus besteht aus fünf Geschäftsberichten des Geschäftsjahres 2023 und fünf Geschäftsberichten des Geschäftsjahres 2024, die von den Webseiten deutscher Unternehmen verschiedener Branchen (vgl. Tab. 1) heruntergeladen wurden.

Tab. 1: Korpusübersicht (Geschäftsberichte 2023–2024) und am Übersetzungsprojekt beteiligte Studierende² (letzter Datenzugriff: 30.10.2025)

Geschäfts-jahr	Unternehmen	Branchen	Anzahl Texte	Studierende (Anzahl / GER-Niveau) und Alter
2024	Deutsche Börse Group, Volkswagen Group, adidas Group, ZF Friedrichshafen AG, DZ Bank AG	Finanzdienstleistungen, Automobilindustrie, Konsumgüter / Einzelhandel, Automobilzulieferer, Bankwesen	5	20 – C1: 18, C2: 2 22–25 Jahre
2023	Beiersdorf AG, Lindt & Sprüngli AG, Rewe Group, Hugo Boss AG, Mercedes-Benz Group AG	Konsumgüterindustrie, Lebensmittel, Handel / Dienstleistungen, Mode, Automobilindustrie	5	25 – C1: 22, C2: 3 22–25 Jahre

.....

2 Eine erweiterte Version der Tabelle mit detaillierten Angaben zu den analysierten Texten ist auf folgendem Padlet abrufbar: <https://padlet.com/busch0lauer/sammelband-wenn-ki-auf-fach-und-sprache-trifft-2026-4dzgd5z5aksxwmxm>.

Aus den Texten wurden Fragmente mit einem Umfang von ca. 2 000 bis 3 000 Wörtern extrahiert, die im Rahmen des Masterkurses *German Specialised Communication and Translation* im Studiengang *Language for Communication in International Enterprise and Organisation (LACOM)* an der Universität Modena und Reggio Emilia in den Jahren 2023/2024 und 2024/2025 übersetzt wurden.

Die Wahl fiel aufgrund ihrer hybriden Struktur auf die Textsorte Geschäftsbericht: Der Geschäftsbericht stellt eine Form der Unternehmenskommunikation dar, die sowohl Informations- als auch Rechenschaftsfunktionen erfüllt; er gibt Auskunft über die wirtschaftliche Lage, die erzielten Ergebnisse sowie die strategische Ausrichtung eines Unternehmens im jeweiligen Geschäftsjahr und leistet damit einen Beitrag zur Imagebildung (Garzone 2004: 311–315; Keller 2006: 15; Hyland 1998: 224–230.). Der Geschäftsbericht vereint in sich verschiedene Subgenres, darunter etwa den Brief an die Aktionäre und den Nachhaltigkeitsbericht, die jeweils durch spezifische kommunikative Zielsetzungen und stilistische Eigenheiten gekennzeichnet sind. Von einem sprachlichen Standpunkt aus ist ein Nominalstil dominierend, der die Objektivität der Information betont. Evaluative Passagen tragen vor allem im Brief an die Aktionäre zur positiven Selbstdarstellung des Unternehmens bei. Somit stellt der Geschäftsbericht eine komplexe, multimodale und strategisch ausgerichtete Textsorte dar, die zwischen Fachkommunikation und öffentlicher Darstellung vermittelt.

Für die Analyse und Übersetzung wurden aus jedem Geschäftsbericht des Korpus die folgenden Textpassagen ausgewählt:

1. Aktionärsbrief/Vorwort des Vorstandsvorsitzenden: Es geht um Texte, die eine dominierende performative Funktion haben. Sie schaffen Vertrauen und vermitteln die Unternehmensstrategie.
2. Lagebericht: Diese Textpassagen informieren über die wirtschaftliche Entwicklung des jeweiligen Unternehmens und sind von einer hohen terminologischen Dichte geprägt.
3. Nachhaltigkeitsberichte: Diese Texte weisen spezifische diskursive Konventionen auf und verorten das jeweilige Unternehmen in der aktuellen Debatte über Nachhaltigkeit.

Diese drei Teile von Geschäftsberichten zeigen terminologische Komplexität, sind kulturspezifisch geprägt (wie z. B. Mittelstand, Standortvorteil, Stakeholder-Dialog) und weisen eine metaphorische Dichte auf. Ferner bieten sie gute Beispiele für typische Übersetzungsschwierigkeiten im Sprachenpaar Deutsch–Italienisch, z. B. Komposita, Nominalstil oder institutionell geprägte Ausdrucksweisen.

5 Zur Analyse der Übersetzungen

Nachdem die Student:innen die Übersetzung ausgewählter Textpassagen selbstständig und nur mit traditionellen Nachschlagewerken vollzogen hatten, übersetzten sie dieselben Textauszüge mit DeepL. Dies erfolgte aus zwei praktischen Gründen: DeepL ist den Student:innen gut bekannt und eignet sich für die Übersetzungsanalyse aufgrund der Möglichkeit, zwischen mehreren lexikalischen Optionen zu wählen (vgl. Maaß/Fioravanti 2024: 475–477; Pielka/Hahnbück/Deußner 2025). Im Gegensatz zu ChatGPT oder Claude führt es keinen Dialog mit dem Übersetzer, sondern liefert direkt den Text, der mit der DeepL-Write-Funktion in einem weiteren Schritt bearbeitet und verbessert werden kann. Dazu ist zu bemerken, dass DeepL im Bereich der fachlichen Übersetzung sowohl in der beruflichen Praxis als auch in der Ausbildung von Übersetzer:innen zunehmend eingesetzt wird, weil es flüssige, formal korrekte Übersetzungen bietet, die oft auf den ersten Blick überzeugend wirken.

In diesem Beitrag werden zur Veranschaulichung ausgewählte Passagen aus dem Aktionärsbrief der DZ Bank³ (Wörterzahl 1965) im Geschäftsbericht 2024 als Beispielmateriale herangezogen. Die ausgewählten Textpassagen sind als Modelltext zu verstehen, weil sie die charakteristischen sprachlichen und stilistischen Besonderheiten eines Geschäftsberichts exemplarisch ver-

.....
 3 Vgl. <https://www.geschaeftsbericht.dzbank.de/de/dz-bank-gruppe/brief-an-die-aktionäre>, letzter Zugriff: 30.10.2025. Da die Zitate alle aus demselben Text stammen, wird die Quelle für jedes Zitat nicht angegeben.

deutlichen und zugleich spezifische Übersetzungsprobleme im Sprachenpaar Deutsch–Italienisch aufzeigen.

Aufbauend auf den ausgewählten Passagen des Aktionärsbriefs erfolgt nun die Analyse, bei der sowohl die studentischen Übersetzungen als auch die maschinellen Übersetzungen von DeepL betrachtet werden.

Textbeispiel (1) zeigt ein typisches Merkmal wirtschaftssprachlicher Fachtexte: Auffällig ist der Nominalstil mit zahlreichen Komposita (*Marktbearbeitung, Genossenschaftsbanken, Kundengeschäft, Konzerngesellschaften*), der die Informationsdichte erhöht und die Objektivität des Textes unterstreicht:

- (1) *Dieses Resultat fußt maßgeblich auf der eng abgestimmten Marktbearbeitung mit den Genossenschaftsbanken und unserem sehr guten Kundengeschäft in den Konzerngesellschaften. Wir konnten so den Wachstumspfad fortsetzen. Die positive Entwicklung an den Kapitalmärkten gab uns zusätzlichen Rückenwind.*

Evaluative Elemente und performative Metaphern (*Rückenwind*) werden hier verwendet, um positive Entwicklungen anschaulich zu machen und motivierend zu vermitteln. Die Sätze variieren zwischen komplexen, informationsreichen Strukturen und kurzen, klaren Folgesätzen, die Ursache und Wirkung logisch verbinden (*Wir konnten so den Wachstumspfad fortsetzen*). Diese sprachlichen Eigenschaften zeigen exemplarisch die Herausforderungen für die Übersetzung ins Italienische: Nominalstil, fachsprachliche Terminologie und metaphorische Ausdrucksweisen erfordern sorgfältige Anpassungen, um sowohl die semantische Präzision als auch die performative Wirkung im Zieltext zu erhalten. Fachtermini wie *Marktbearbeitung, Genossenschaftsbanken, Kundengeschäft, Konzerngesellschaften, Wachstumspfad* und *Kapitalmärkte* können für die Student:innen besonders schwierig sein. Die kollektive Übersetzung der Student:innen im Klassenzimmer lautet wie folgt:

- (1a) *Questo risultato si basa soprattutto sullo sviluppo del mercato strettamente concordato con le banche di credito cooperativo e il nostro ottimo lavoro con i clienti delle società consorziate. Abbiamo potuto*

così proseguire il nostro percorso di crescita. Lo sviluppo positivo dei mercati dei capitali ci ha dato ulteriore slancio.

Die Übersetzung ist insgesamt korrekt und gut verständlich. Der Satz *Abbiamo potuto così proseguire il nostro percorso di crescita* entspricht den italienischen Satzbaugewohnheiten und wirkt idiomatisch. Auch die Terminologie ist präzise ausgewählt worden: *mercati dei capitali* sowie *ulteriore slancio* reflektieren den formalen wirtschaftssprachlichen Stil des Originals. Stilistisch gelingt es, sowohl die informative als auch die evaluative Funktion zu bewahren, da der Text Vertrauen signalisiert und die erfolgreiche Unternehmensentwicklung hervorhebt. Die Student:innen haben die Metapher *Rückenwind*, die den zusätzlichen Antrieb durch positive Marktentwicklungen bildhaft vermittelt, nicht übersetzt, sondern den Inhalt mit einem neutralen Ausdruck wiedergegeben. Formulierungen wie *il nostro ottimo lavoro con i clienti* sind verständlich, stilistisch aber etwas freier wiedergegeben. Besonders anspruchsvoll waren Ausdrücke wie *fußt maßgeblich auf*, *eng abgestimmte Marktbearbeitung*, *Wachstumspfad fortsetzen* und Komposita wie *Genossenschaftsbanken* und *Konzerngesellschaften*. Bei *fußt maßgeblich auf* mussten die Studierenden sowohl die sachliche Grundlage der Aussage als auch die evaluative und performative Komponente sowie das verstärkende Adverb *maßgeblich* angemessen ins Italienische übertragen. *Eng abgestimmte Marktbearbeitung* erforderte die präzise Wiedergabe der Fachterminologie und gleichzeitig die Erfassung der Nuance der engen Kooperation. Begriffe wie *Genossenschaftsbanken* und *Konzerngesellschaften* stellten besondere Herausforderungen dar, da sie stark an das deutsche Bankensystem bzw. an spezifische Unternehmensstrukturen gebunden sind und nicht wortwörtlich übersetzt werden können. Solche Unterschiede stellen typische Herausforderungen der Fachübersetzung wirtschaftlicher Texte dar: Neben semantischer Genauigkeit müssen Übersetzer:innen auch die pragmatische Wirkung, den stilistischen Registercharakter und metaphorische Elemente berücksichtigen, um die volle Funktion des Originaltexts zu erhalten.

Die Übersetzung von DeepL ist wie folgt (1b):

- (1b) *Questo risultato è dovuto in gran parte allo sviluppo del mercato strettamente coordinato con le banche cooperative e all'ottima attività con i clienti delle società del Gruppo. Siamo quindi riusciti a proseguire il nostro percorso di crescita. L'andamento positivo dei mercati dei capitali ci ha dato un'ulteriore spinta.*

Die DeepL-Übersetzung stellt sich als flüssig, idiomatisch und professionell dar, weist jedoch einige Einschränkungen auf. Fachtermini wie *Genossenschaftsbanken* werden mit *banche cooperative* übersetzt, was zwar verständlich ist, jedoch die spezifische institutionelle Bedeutung im deutschen Bankensystem nur annähernd wiedergibt. Auch *Konzerngesellschaften* wird durch *società del Gruppo* verallgemeinert, wodurch die komplexe Struktur der Unternehmensgruppen abgeschwächt wird. Darüber hinaus wirken einige evaluative und performative Nuancen, wie sie etwa in *maßgeblich auf* oder *Rückenwind* enthalten sind, in der italienischen Version leicht neutralisiert. Schließlich ist das formale Register etwas weniger präzise und formell als im Original, etwa bei *Siamo quindi riusciti a proseguire il nostro percorso di crescita*. Insgesamt zeigt sich, dass DeepL zwar eine stilistisch elegante und verständliche Übersetzung liefert, aber in Bezug auf Fachpräzision und performative Wirkung noch Optimierungspotenzial besteht.

In Textbeispiel (2) ist das Verb *belasten* ein Euphemismus, der wirtschaftliche Schwierigkeiten beschreibt. Es verleiht dem Text ein neutrales Register und vermeidet, der Situation einen dramatischen Ton zu geben:

- (2) *Die schwache Konjunkturlage und die damit verbundenen erhöhten Ausfallrisiken belasteten einzelne Institute der DZ BANK Gruppe.*

In diesem Beispiel zeigt sich ein deutlicher Unterschied im sprachlichen Register und in der semantischen Nuancierung zwischen der studentischen (2a) und der maschinellen Übersetzung (2b):

- (2a) *La debole congiuntura economica e gli aumentati rischi di insolvenza a esse collegata hanno penalizzato singoli istituti del gruppo DZ Bank.*

- (2b) *La debolezza della situazione economica e il conseguente aumento del rischio di insolvenza hanno avuto un impatto negativo sui singoli istituti del Gruppo DZ BANK.*

Die studentische Übersetzung *hanno penalizzato* (2a) wirkt stärker wertend und emotional gefärbt, da *penalizzare* eine aktiv verursachte Benachteiligung suggeriert. Die DeepL-Version *hanno avuto un impatto negativo* (2b) entspricht dem neutralen Charakter des Ausgangsverbs präziser. Sie bewahrt auch den objektiven, analytischen Stil, der für Geschäftsberichte typisch ist. Deswegen passt diese Lösung besser zum Register des Originals, auch wenn die Formulierung etwas länger und weniger komprimiert wirkt als dt. *belasten*. Trotzdem weist die DeepL-Übersetzung einige Einschränkungen auf. Die Formulierung *hanno avuto un impatto negativo* spiegelt zwar den neutralen Ton des Verbs *belasten* relativ gut wider, wirkt jedoch etwas schwerfällig und formelhaft. Zudem verliert die Übersetzung die implizite, kontinuierliche Wirkung, die *belasten* im Deutschen vermittelt, da das italienische Ausdrucksmittel ein punktuell Ereignis suggeriert. Der Ausdruck klingt daher leicht distanziert und weniger dynamisch als das Original.

In Textbeispiel (3) sehen sich die Student:innen mit einem Ausdruck konfrontiert, der fachliche Kompetenzen voraussetzt. Dazu stellt die Kombination aus einer quantitativen Angabe (*15,8 Prozent*) und einer qualitativen Bewertung (*sehr robust*) für die Übersetzung eine besondere Herausforderung dar:

- (3) *Mit einer harten Kernkapitalquote von 15,8 Prozent verfügt die DZ BANK Gruppe weiterhin über eine sehr robuste Kapitalausstattung.*

Im Deutschen wird durch diese Struktur eine doppelte Betonung erzeugt: eine objektive, messbare Stärke und eine subjektive, positive Bewertung. Bei der Übertragung ins Italienische besteht die Gefahr, dass die semantische Balance verloren geht. Eine wörtliche Übersetzung wie *molto solida* bewahrt zwar die positive Konnotation, schwächt jedoch den bewertenden Charakter von *sehr robust*, was im Deutschen zugleich Stabilität, Widerstandsfähigkeit und Zuverlässigkeit impliziert. Außerdem kann die quantitative Zahl im Italienischen weniger stark als Argument der Glaubwürdigkeit wirken, da der Diskurs im

Zielkontext häufig stärker narrativ und weniger technisch argumentativ geprägt ist. Übersetzer:innen müssen daher eine Lösung finden, die sowohl die faktische Dimension als auch die evaluative Intention beibehält.

Die studentische Übersetzung lautet:

- (3a) *Con un solido coefficiente patrimoniale di base del 15,8 % il gruppo DZ Bank possiede un consistente capitale.*

Diese Übersetzung gibt *harte Kernkapitalquote* mit *coefficiente patrimoniale di base* wieder, was zwar verständlich ist, aber nicht der offiziellen Fachterminologie des Bank- und Finanzwesens entspricht. Im internationalen und italienischen Kontext wird hierfür üblicherweise der englische Ausdruck *Common Equity Tier 1 (CET1)* verwendet, der auch in offiziellen Dokumenten der Banca d'Italia oder der Europäischen Zentralbank vorkommt.

DeepL (3b) wählt genau diese Lösung und zeigt damit eine terminologisch präzisere und an den institutionellen Diskurs angepasste Übersetzungsstrategie:

- (3b) *Con un coefficiente patrimoniale Common Equity Tier 1 del 15,8 %, il Gruppo DZ BANK continua a disporre di una base patrimoniale molto solida.*

Ferner ist der Satzbau in der studentischen Version relativ schlicht und wirkt etwas holprig, insbesondere durch die Formulierung *possiede un consistente capitale*, die die Bedeutung von *Kapitalausstattung* zu eng fasst und eher auf das Eigentum an Kapital verweist. Dagegen verwendet DeepL *disporre di una base patrimoniale molto solida*, was mit dem formalen Register von Geschäftsberichten idiomatischer und stilistisch kohärent ist. Positiv hervorzuheben ist auch, dass DeepL das Adverb *weiterhin* korrekt in der Formulierung *continua a disporre* wiedergibt, wodurch die Kontinuität der stabilen Kapitalstruktur sprachlich erhalten bleibt. Obwohl die DeepL-Übersetzung überzeugender wirkt, zeigt sie eine gewisse Distanz zum Ausgangstext: Der Ausdruck *base patrimoniale molto solida* wirkt generalisierend und verliert die spezifische technische Bedeutung von *Kapitalausstattung*.

Textbeispiel (4) enthält eine Partizipialkonstruktion mit informativer Dichte:

- (4) *Unterstützt durch die positive Entwicklung an den Aktienmärkten erzielten sowohl die R+V Versicherung als auch die Union Investment sehr gute Ergebnisse.*

Sowohl die studentische Übersetzung (4a) als auch die von DeepL (4b) geben den Inhalt des Originalsatzes korrekt wieder, unterscheiden sich jedoch in Nuancen der Syntax, Idiomatik und stilistischen Präzision:

- (4a) *Sostenute dallo sviluppo positivo dei mercati azionari sia l'assicurazione R+V sia l'Union Investment hanno raggiunto ottimi risultati.*
- (4b) *Sostenute dall'andamento positivo dei mercati azionari, sia R+V Versicherung che Union Investment hanno ottenuto ottimi risultati.*

In der studentischen Version wird die Satzstruktur des Deutschen beibehalten, obwohl der Einstieg *Sostenute dallo sviluppo positivo dei mercati azionari* etwas schwerfällig wirkt. Im Italienischen würde ein Komma nach dem Partizipialkonstrukt die Lesbarkeit verbessern. Auch die Reihenfolge von Subjekt und Verb (*sia l'assicurazione R+V sia l'Union Investment hanno raggiunto*) wirkt leicht unbalanciert, da die Apposition *l'assicurazione* hier eher erläuternden Charakter hat und stilistisch in institutionellen Texten oft ausgelassen wird. Die DeepL-Version zeigt hingegen eine flüssigere Syntax und bessere idiomatische Wortwahl: *andamento positivo* ist im Finanzkontext üblicher als *sviluppo positivo* und entspricht den dynamischen Marktbewegungen. Die Übersetzung *hanno ottenuto ottimi risultati* ist präzise und stilistisch angemessen.

Insgesamt nähert sich die maschinelle Übersetzung stärker der italienischen Wirtschaftssprache an, während die studentische Version eine gewisse formale Schwerfälligkeit beibehält. Beide Varianten verzichten auf die leichte evaluative Färbung des deutschen Originals *sehr gute Ergebnisse*, die im Italienischen durch *risultati molto positivi* etwas neutraler wirkt. Das Beispiel verdeutlicht, wie kleine lexikalische und syntaktische Entscheidungen den Grad

der Fachangemessenheit und den institutionellen Ton eines Wirtschaftstextes beeinflussen.

Das folgende Textbeispiel (5) ist durch eine hybride fachlich-bildliche Terminologie charakterisiert:

- (5) *Es ist daher zu hoffen, dass zügig eine neue Bundesregierung gebildet wird, die die drängenden Strukturreformen zur Sicherung von Wachstum, Wohlstand, Sicherheit und Infrastruktur in unserem Land unmittelbar angeht.*

Die Formulierung *drängende Strukturreformen* verbindet das Bild einer physischen Erpressung mit einem ökonomischen Terminus. *Unmittelbar angehen* wirkt umgangssprachlich in einem formellen Kontext. Die Wortreihe *Wachstum, Wohlstand, Sicherheit und Infrastruktur* hat einen alliterativen Rhythmus (zwei W-) mit präskriptiver Kraft, der in der italienischen Übertragung verloren geht. Auch die Formulierungen *es ist zu hoffen, unmittelbar angeht* haben einen appellativen Charakter: Die deutsche Sprache bewegt sich zwischen Wunsch und moralischer Verpflichtung. In ihrer Übersetzung zeigen die Student:innen, dass sie den Sinn richtig verstanden haben, jedoch bleiben sie noch dem Ausgangstext zu nah:

- (5a) *Si spera quindi che venga presto formato un nuovo governo che si occupi subito delle pressanti riforme strutturali in materia di crescita, benessere, sicurezza e infrastrutture nel nostro paese.*

In der studentischen Übersetzung *Si spera quindi che venga presto formato un nuovo governo ...* wird der unpersönliche und appellative Charakter des Originals beibehalten, jedoch klingt *si spera* schwächer und weniger institutionell.

Die DeepL-Übersetzung (5b) bietet mit *è quindi auspicabile* eine stilistisch gehobenere und im politischen Diskurs angemessenere Lösung. Auch *affrontare le urgenti riforme strutturali* trifft den formellen und auffordernden Ton des Originals besser als *occuparsi subito*, das alltagssprachlicher wirkt. Jedoch verliert auch DeepL die bildhafte Komponente von *drängend* und die rhythmische Wirkung der deutschen Wortreihe:

- (5b) *È quindi auspicabile che si formi rapidamente un nuovo governo federale per affrontare le urgenti riforme strutturali necessarie a garantire la crescita, la prosperità, la sicurezza e le infrastrutture del nostro Paese.*

Dieses Beispiel zeigt besonders deutlich, wie stark semantische Feinheiten, rhetorische Strukturen und Registerentscheidungen den Übersetzungsprozess wirtschaftspolitischer Texte prägen. Während DeepL eine formell angemessene und idiomatische Lösung anbietet, gelingt es weder der maschinellen noch der studentischen Übersetzung vollständig, die metaphorische Spannung und den appellativen Rhythmus des deutschen Originals zu bewahren.

Im folgenden Satz (6) sind die Student:innen mit drei Elementen konfrontiert: einem komplexen Subjekt (*die Verbesserung der Wettbewerbsfähigkeit*); einem explikativen Einschub (*auch durch den Abbau*), und einem idiomatischen Prädikat (*duldet keinen weiteren Aufschub*):

- (6) *Insbesondere die Verbesserung der Wettbewerbsfähigkeit der deutschen Wirtschaft, auch durch den Abbau bürokratischer Hemmnisse, duldet keinen weiteren Aufschub.*

Die für das Verwaltungsdeutsch typische Anhäufung von Genitiven und die verpflichtende politische Sprache können zu Problemen bei der Übersetzung führen. *Abbau* ist eine konstruktive Metapher, die auf einen allmählichen Prozess hinweist, nicht auf einen sofortigen. Die Verwendung von *auch* in der Einfügung deutet an, dass der Abbau von Bürokratie eine der möglichen Lösungen ist. Beide Übersetzungen, die studentische (6a) und die mit DeepL (6b) zeigen ein gutes Verständnis des Ausgangstextes, unterscheiden sich aber in Stil, Idiomatik und Grad der Zieltextadäquatheit:

- (6a) *Soprattutto il miglioramento della capacità concorrenziale dell'economia tedesca, anche grazie allo smantellamento degli intoppi burocratici, non può tollerare nessun ulteriore ritardo.*

- (6b) *In particolare, il miglioramento della competitività dell'economia tedesca, anche attraverso la riduzione degli ostacoli burocratici, non può essere più rimandato.*

In der studentischen Version fällt die wörtliche Struktur auf: *soprattutto il miglioramento della capacità concorrenziale dell'economia tedesca* ist grammatikalisch korrekt, jedoch weniger idiomatisch. Der Ausdruck *capacità concorrenziale* stellt eine semantisch plausible, jedoch zu technische Entsprechung von *Wettbewerbsfähigkeit* dar; im Fach- und Verwaltungskontext ist der Terminus *competitività* die konventionellere und stilistisch passendere Lösung. Auch die Übersetzung von *Abbau* mit *smantellamento* bleibt bildlich nah am Deutschen, wirkt jedoch im institutionellen Italienisch etwas zu stark und aktivistisch, während *riduzione*, wie in (6b), den graduellen, prozesshaften Charakter des Begriffs besser wiedergibt.

Was die Satzorganisation betrifft, klingt die studentische Version durch die fehlende Einleitungspartikel *In particolare* weniger rhetorisch. (6b) verwendet dagegen eine flüssigere Struktur, die dem italienischen Diskursstil entspricht, und integriert die adverbiale Bestimmung elegant. Besonders gelungen ist in (6b) auch die Wiedergabe des idiomatischen Prädikats *duldet keinen weiteren Aufschub*. Während (6a) mit *non può tollerare nessun ulteriore ritardo* eine wörtlichere, aber leicht unidiomatische Variante bietet, wird in (6b) mit *non può essere più rimandato* der performative, politische Ton des Originals präziser und natürlicher wiedergegeben. Ferner weist (6b) einen höheren Grad an Idiomatik, stilistischer Kohärenz und funktionaler Äquivalenz auf. Sie wahrt den formellen, politischen Ton des Verwaltungsdeutschen, ohne die Lesbarkeit im Italienischen zu beeinträchtigen.

In der folgenden Textpassage (7) geht es sowohl um ein terminologisches als auch um ein metaphorisches Problem:

- (7) *Deutschland hat weiterhin eine starke Substanz, sei es in der Grundlagenforschung oder in der Ausbildung von Fachkräften, auf der Gesellschaft und Wirtschaft aufbauen können.*

Das Wort *Substanz* wird hier im metaphorischen Sinn verwendet, darf deshalb nicht mit *sostanza*, sondern mit *sostrato*, *base solida*, *fondamenta* übersetzt werden. *Substanz haben* ist ein idiomatischer Ausdruck, der aber kein direktes Äquivalent auf Italienisch hat:

Die studentische Version (7a) verwendet *solide fondamenta*, übersetzt aber den Relativsatz unidiomatisch (*sia che si tratti ...*), wodurch die Lesbarkeit leidet und die Syntax vom Italienischen abweicht. Auch die Verbindung zwischen den einzelnen Elementen (*ricerca di base*, *formazione di personale specializzato*, *società e l'economia*) ist nicht optimal integriert:

(7a) *La Germania ha inoltre solide fondamenta, sia che si tratti della ricerca di base sia nella formazione di personale specializzato sul quale la società e l'economia possono crescere.*

Die DeepL-Version (7b) zeigt eine idiomatischere und flüssigere Wiedergabe: *continua ad avere una solida base* überträgt das metaphorische Konzept von *Substanz haben* treffend. Der Relativsatz *su cui la società e l'economia possono costruire* entspricht der logischen Struktur des Ausgangstextes, bewahrt somit die metaphorische Bedeutung, dass Forschung und Ausbildung die Grundlage für gesellschaftliche und wirtschaftliche Entwicklung bilden:

(7b) *La Germania continua ad avere una solida base, sia nella ricerca di base che nella formazione di manodopera qualificata, su cui la società e l'economia possono costruire.*

Das letzte Beispiel (8) illustriert die Bedeutung feiner semantischer und pragmatischer Nuancen in der Übersetzung wirtschaftlicher Fachtexte. Das deutsche Verb *gestalten* impliziert aktives, schöpferisches Handeln und sollte daher im Italienischen mit einem Verb wiedergegeben werden, das die kreative Dimension betont (*plasmare*, *forgiare* etc.):

(8) *Die Überzeugung und das Vertrauen, unsere Zukunft selbst gestalten zu können, sollten uns bei den vor uns liegenden Aufgaben leiten.*

In der studentischen Version (8a) wurde *creare* gewählt, was verständlich ist, aber weniger idiomatisch wirkt und die schöpferische Dimension des Originals nicht vollständig wiedergibt:

(8a) *La condizione e la fiducia di poter noi stessi creare il nostro futuro dovrebbero guidarci nelle sfide che dobbiamo affrontare.*

Die DeepL-Version (8b) verwendet das Verb *plasmare*, das sowohl die aktive als auch die gestaltende Komponente aufnimmt und die metaphorische Intention des Originals reflektiert:

(8b) *La convinzione e la fiducia di poter plasmare il nostro futuro dovrebbero guidarci nei compiti che ci attendono.*

Beide Übersetzungen vermitteln die wechselseitige Verbindung von Überzeugung, Vertrauen und Wunschcharakter der Aussage. Dabei ist wichtig zu beachten, dass der Konjunktiv II im Deutschen weder eine starke Verpflichtung noch einen simplen Ratschlag ausdrückt, sondern einen Wunsch oder eine motivierende Aufforderung signalisiert. In beiden italienischen Versionen wird die Modalität des Konjunktivs II, d. h., einen Wunsch oder eine motivierende Aufforderung zu zeigen, durch die Konstruktion *dovrebbero guidarci* übertragen.

Die Analyse der studentischen Übersetzungen erlaubt nicht nur die Einordnung typischer Übersetzungsprobleme, sondern auch die Reflexion über den Lernprozess. Diese Übersetzungen beweisen, dass die Auseinandersetzung mit terminologischer Präzision und den kulturellen Konnotationen wirtschaftlicher Schlüsselbegriffe zu einem vertieften Bewusstsein für die Diskursgebundenheit ökonomischer Sprache führt. Dies manifestiert sich insbesondere in der Entscheidung der Studierenden zwischen formal korrekten, jedoch stilistisch unpassenden Lösungen und kontextuell adäquateren Varianten. Der Vergleich mit den DeepL-Übersetzungen intensivierte den Reflexionsprozess: Einerseits erfassen maschinelle Systeme die semantische Dimension korrekt, andererseits neigen sie dazu, pragmatische und diskursive Funktionen zu neutralisieren. Die Studierenden erkennen dadurch den Mehrwert menschlicher

Entscheidungskompetenz, die Bedeutungsnuancen zu erkennen und kommunikative Intentionen zu vermitteln. Darüber hinaus wurde deutlich, dass der Einsatz von KI im Unterricht als Katalysator metakognitiver Prozesse wirkt. Dadurch können die Studierenden ein kritisches Bewusstsein für Übersetzungsentscheidungen entwickeln, das über das bloße Korrigieren maschineller Vorschläge hinausgeht und ihre eigene translatorische Verantwortung stärkt.

6 Vorläufiges Fazit

Die vorliegende Untersuchung hat gezeigt, dass KI-gestützte Systeme wie DeepL wirtschaftliche Fachtexte flüssig, idiomatisch und formal kohärent ins Italienische übertragen können. Insbesondere bei der Terminologie, der Satzstruktur und der Lesbarkeit erreichen maschinelle Übersetzungen ein hohes Niveau und bieten Studierenden sowie Übersetzer:innen wertvolle Orientierungshilfen. Gleichzeitig verdeutlichen die Beispiele (1–8), dass bestimmte Aspekte menschlicher Übersetzungsleistung nach wie vor unersetzlich bleiben: Die präzise Wiedergabe metaphorischer Elemente, die feine Abstimmung zwischen pragmatischer Intention und stilistischer Registertreue sowie die semantische Differenzierung kultur- und fachspezifischer Begriffe erfordern ein tiefgehendes Verständnis des Ausgangstextes, das über die Fähigkeiten der KI hinausgeht.

Die Übersetzung wirtschaftlicher Texte benötigt nicht einfach ein sprachliches Handeln, sondern auch ein kognitiv-strategisches. Während DeepL eine solide Grundlage liefert und die Effizienz des Übersetzungsprozesses erhöhen kann, bleibt die aktive Rolle des menschlichen Übersetzers entscheidend, um sowohl die kommunikative Wirkung als auch die performative Dimension der Texte zu bewahren.

Zur Beantwortung der einleitenden Forschungsfragen lässt sich feststellen, dass KI-gestützte Systeme wie DeepL Studierenden eine äußerst hilfreiche Unterstützung beim Übersetzen bieten: Sie erzeugen idiomatische Satzstrukturen, gewährleisten terminologische Konsistenz und liefern einen ersten Entwurf, der als Grundlage für die weitere Überarbeitung dienen kann. Auf diese

Weise können Studierende lernen, Unterschiede zwischen maschineller und menschlicher Übersetzung zu erkennen und gezielt Optimierungen vorzunehmen – etwa bei der metaphorischen Übertragung oder der pragmatischen Feinabstimmung. Metaphern wie *Rückenwind* oder idiomatische Wendungen wie *duldet keinen weiteren Aufschub* werden von DeepL verständlich übersetzt, verlieren jedoch ihre evaluative, performative oder metaphorische Nuance. Dazu zeigt sich auch, dass institutionell spezifische Termini (z. B. *Genossenschaftsbanken*, *Konzerngesellschaften*) verallgemeinert werden, was die Fachpräzision einschränkt. Die studentischen Übersetzungen verdeutlichen dagegen, dass menschliche Übersetzer:innen oft unbewusst auf pragmatische und stilistische Anpassungen zurückgreifen, die KI-Systeme nicht leisten, wie die differenzierte Wiedergabe von euphemistischen Verben, Appellformen oder aktiven, schöpferischen Handlungen (vgl. „bebelasten“). Beide Übersetzungen der Beispiele haben bewiesen, dass sich die Grenzen dieser Systeme vorwiegend in der semantischen Präzision, beim Umgang mit Metaphern, idiomatischen Ausdrücken und kulturell spezifischen Begriffen zeigen. KI tendiert dazu, performative und evaluative Nuancen zu neutralisieren bzw. institutionell spezifische Termini zu verallgemeinern. Die Student:innen darauf aufmerksam zu machen, bedeutet vor allem, sie für die Analyse von Registern, Diskursfunktionen und kulturellen Besonderheiten zu sensibilisieren.

Im Bewusstsein, dass dieser Aufsatz und die zugrunde liegende Studie lediglich einen begrenzten Beitrag zur Erforschung der Rolle von KI im fachlichen Kontext leisten, eröffnen sie dennoch Denkanstöße für potenzielle zukünftige Forschungs- und Entwicklungsperspektiven: Wird KI im Unterricht als Werkzeug eingesetzt, das automatisierte Analysen mit qualitativer Reflexion verbindet, kann sie zur Förderung didaktischer und linguistischer Kompetenzen beitragen. Lernen Studierende beispielsweise, systematisch Unterschiede zwischen maschineller und menschlicher Übersetzung zu analysieren und zu kommentieren, sind sie besser in der Lage, die fortlaufende Entwicklung von KI-Modellen zu verstehen, sie kritisch und gezielt zu nutzen – etwa im Hinblick auf kulturelle, metaphorische und pragmatische Nuancen. Es erscheint durchaus denkbar, dass sich hybride Ansätze, bei denen maschinelle Vorschläge durch menschliche Nachbearbeitung ergänzt werden, langfristig als Standardpraxis in der Fachübersetzung etablieren.

Literatur

- BIEL, ŁUCJA / SOSONI, VILELMINI (2017): „The Translation of Economics and the Economics of Translation.“ *Perspectives* 25.3: 351–361. DOI: 10.1080/0907676X.2017.1313281.
- DEUTSCHER GENOSSENSCHAFTS- UND RAIFFEISENVERBAND E. V. (DGRV) (2006): *Studie: Die Abgrenzung von Eigen- und Fremdkapital nach IFRS. Ein Vorschlag zur Modifizierung von IAS 32.* (PDF). <https://www.dgrv.de/publication/die-abgrenzung-von-eigenkapital-und-fremdkapital-nach-ifs/> (31.10.2025).
- FUERTES-OLIVERA, PEDRO A. / NIELSEN, SANDRO (2012): „Online Dictionaries for Assisting Translators of LSP Texts: The Accounting Dictionaries.“ *International Journal of Lexicography* 25.2: 191–215. DOI:10.1093/ijl/ecr034.
- GARZONE, GIULIANA ELENA (2004): „Annual Company Reports and CEO’s Letters: Discoursal Features and Cultural Markedness.“ *Intercultural Aspects of Specialized Communication*. Hrsg. Christopher N. Candlin / Maurizio Gotti. Bern: Lang: 311–342.
- GERHALTER, KATHARINA (2024): „How do DeepL and ChatGPT process information structure and pragmatics? An exploratory case study on topicalized infinitives in Spanish (and Portuguese).“ *AI-Linguistica. Linguistic Studies on AI-Generated Texts and Discourses* vol. 1: 1–36. DOI: 10.62408/ai-ling.v1i1.8.
- HIEBL, BETTINA / GROMANN, DAGMAR (2023): „Quality in Human and Machine Translation: An Interdisciplinary Survey.“ *Proceedings of the 24th Annual Conference of the European Association for Machine Translation*. Tampere. 375–384.
- HYLAND, KEN (1998): „Exploring Corporate Rhetoric: Metadiscourse in the CEO’s Letter.“ *Journal of Business Communication* 35: 224–245.
- JANISCH, MONIKA (1993): *Das strategische Anspruchsgruppenmanagement: vom Shareholder Value zum Stakeholder*. Bern: Haupt.
- KELLER, RUDI (2006): *Der Geschäftsbericht. Überzeugende Unternehmenskommunikation durch klare Sprache und gutes Deutsch*. Wiesbaden: Gabler.
- KIRALY, DONALD C. (2014): *A Social Constructivist Approach to Translator Education*. London: Taylor & Francis.
- LEMON, TETYANA / FLUCK, HANS-RÜDIGER (2016): *Kognitive Metaphernmodelle in der Wirtschaftskommunikation: Eine Empirische Studie zu Geschäftsberichten Deutscher und Russischer Banken*. Berlin: Stauffenburg.

- LÜGER, HEINZ-HELMUT (2018): „Äquivalenz und kommunikative Gleichwertigkeit – Übersetzungsprobleme im sprachwissenschaftlichen Fachtext.“ *Fachsprache Französisch – facettenreich, faszinierend, funktional*. Hrsg. Thomas Tinnefeld. Saarbrücken: htw saar. 71–86.
- LUTZ, BENEDIKT (2017): „Modelle für die verständliche Fachkommunikation. Das Spannungsfeld zwischen wissenschaftlichem Anspruch und praktischer Anwendbarkeit.“ *trans-kom* 10.3: 284–314.
- MAASS, CHRISTIANE / FIORAVANTI, CHIARA (2024): „Evaluating the performance of DeepL as translation tool between German and Italian Easy Language administrative texts.“ *Rivista italiana di informatica e diritto*: 475–500. DOI: 10.32091/RIID0222.
- O'BRIEN, SHARON / SIMARD, MICHEL / SPECIA, LUCIA (eds.) (2012): *Proceedings of the First Workshop on Post-Editing Technologies and Practice (WPTP1)*. Association for Machine Translation in the Americas, San Diego, California, USA. <https://aclanthology.org/2012.amta-wptp.0/> (10.10.2025).
- PIELKA, MAREN / HAHNBÜCK, MAX / DEUSSER, TOBIAS et al. (2025): „Automating translation checks of financial documents using large language models.“ *Lang Resources & Evaluation*. DOI: 10.1007/s10579-025-09862-z.
- RICO, CELIA / GONZÁLEZ PASTOR, DIANA (2022): „The Role of Machine Translation in Translation Education: A Thematic Analysis of Translator Educator's Beliefs.“ *Translation and Interpreting: International Journal for Translation and Interpreting Research* 14.1: 177–197.
- RISKU, HANNA / ROGL, REGINA / MILOŠEVIĆ, JELENA (eds.) (2017): „Translation Practice in the Field: Current Research on Socio-cognitive Processes.“ *Special Issue Translation Spaces* 6.1: 3–26.
- ROELCKE, THORSTEN (2014): „Zur Gliederung von Fachsprache und Fachkommunikation.“ *Fachsprache* 36.3–4: 154–178.
- SCARPA, FEDERICA (2020): *Research and Professional Practice in Specialised Translation*. *Palgrave Studies in Translating and Interpreting*. Palgrave-Macmillan.
- SCHROPP, HELEN/DIPPOLD-SCHENK, KATJA (2018): „Integriertes Fach- und Sprachlernen in der beruflichen (Anpassungs-)Qualifizierung.“ *Sprachen im Beruf* 1.2: 135–152.

- SIMONNÆS, INGRID/ KRISTIANSEN, MARITA (2018): „Economic Language.“ *Languages for Special Purposes: An International Handbook*. Hrsg. John Humbley / Gerhard Budin / Christer Laurén. Berlin: De Gruyter. 151–208.
- STOLZE, RADEGUNDIS (2009): *Fachübersetzen – Ein Lehrbuch für Theorie und Praxis*. (Forum für Fachsprachenforschung 89). Berlin: Frank & Timme.
- TAUS (2017): *Machine Translation Post-Editing Guidelines* (TAUS DQF / MTPE Guidelines). <https://www.taus.net/resources/reports/mt-post-editing-guidelines> (15.01.2026).
- TCACIUC, LUCIANA / MACKEVIC, VLADISLAV (2017): „Translators’ Agency in Translating Economic Metaphors in European Union Institutions: the Case of the European Central Bank.“ *Perspectives* 25.3: 417–433.
- WIESMANN, EVA (2018): „Kulturspezifik in Recht und Technik und Konsequenzen für die Übersetzung.“ *inTRAlinea* 20. https://www.intralinea.org/specials/article/kulturspezifik_in_recht_und_technik_und_konsequenzen_fuer_die_uebersetzung (20.11.2025).
- WITCZAK, OLGA (2016): „Incorporating Post-Editing into a Computer-assisted Translation Course: A Study of Student Attitudes.“ *Journal of Translator Education and Translation Studies* 1.1: 33–55.
- YANG, XIN YI / ZHAN, RUNZHE / WONG, DEREK F. / WU, JUN CHAO / CHAO, LIDIA S. (2023): „Human-in-the-loop Machine Translation with Large Language Model.“ *Proceedings of Machine Translation Summit XIX*. <https://arxiv.org/abs/2310.08908> (02.10.2025).
- ZANETTIN, FEDERICO (2018): „Electronic Tools and Resources for Translating and Writing in the Digital Age.“ *inTRAlinea* 20 https://www.intralinea.org/specials/article/electronic_tools_and_resources_for_translating_and_writing_in_the_digital (20.11.2025).
- ZAYKOVA, IRINA / IRINA SHILNIKOVA (2019): „Economic Translation: Theoretical and Practical Issues.“ *SHS Web of Conferences (CILDIH) 2019. The International Scientific and Practical Conference „Current Issues of Linguistics and Didactics: The Interdisciplinary Approach in Humanities and Social Sciences“ (CILDIAH-2019)*. *SHS Web Conf.* 69. DOI: 10.1051/shsconf/20196900139.

RUI MA

When AI Tastes Iberian Ham

Cultural Gaps in Spanish–Chinese Machine Translation of Gourmet Tourism Products

Abstract

This study examines the cultural and discursive gaps that emerge when machine translation systems render Spanish gourmet tourism products into Chinese, focusing particularly on texts involving Iberian ham. While neural machine translation models such as DeepL, ChatGPT, Kimi, DeepSeek and Copilot have made significant progress in lexical accuracy and contextual fluency, they still struggle to convey symbolic, sensory and culturally embedded meanings characteristic of gastronomic discourse. Through a comparative analysis of automatic outputs and translations produced by university students, this paper identifies recurrent patterns of cultural neutralisation, loss of aesthetic tone and weakened promotional discourse in machine-generated versions. By contrast, human translators employ interpretive strategies, such as amplification, domestication and idiomatic recreation, that restore cultural depth and communicative intent. The findings highlight the indispensable role of cultural mediation in translation training and argue for a critical, pedagogical integration of AI to foster intercultural competence and translational creativity in future professionals.

Keywords cultural mediation, gastronomic translation, Iberian ham, machine translation, Spanish–Chinese translation

1 Introduction

The rapid expansion of Artificial Intelligence (AI) across academic and professional domains has profoundly reshaped translation practices. Gastronomic discourse, especially when linked to tourism promotion, represents a particularly demanding field for AI-based translation due to its dense cultural symbolism, sensory imagery and implicit aesthetic values. In the Spanish–Chinese language pair, these challenges become especially pronounced, as food operates as a key site of identity construction, social memory and cultural expression in both traditions.

Products such as Iberian ham and olive oil are not mere commodities within Spanish culture; they are emblematic symbols tied to craftsmanship, heritage and national pride. Their translation into Chinese requires more than terminological equivalence. It involves transmitting a cultural universe shaped by history, rituals and affective associations. While neural machine translation systems have achieved remarkable improvements in morphosyntactic accuracy, they continue to struggle when confronted with culturally encoded meanings, symbolic references or promotional tones that rely on sensory evocation.

In this context, human translators play a decisive role as cultural mediators. Their intervention ensures that the cultural, emotional and symbolic dimensions of gastronomic discourse are preserved in the target text. The present study compares machine translations of Spanish gastronomic descriptions with translations produced by university students in a controlled classroom setting. The aim is to identify not only the recurring shortcomings of machine-generated outputs but also the interpretive, creative and culturally informed strategies employed by human translators.

The analysis sheds light on the limits of automated systems when handling culturally saturated texts, while emphasizing the pedagogical relevance of integrating AI as a reflective tool in translation training. By juxtaposing both approaches, this study highlights the irreplaceable value of human cultural competence and demonstrates how translation pedagogy can benefit from a balanced, critical use of AI technologies.

2 Theoretical frameworks

2.1 Gastronomy as cultural and symbolic discourse

Food practices are deeply embedded in the social and historical life of communities. As Carretero (2014) notes, eating is not merely a biological act but a cultural performance that shapes identities and mediates interpersonal relations. In Spain, products such as Iberian ham or olive oil occupy a prominent place in national imagination. They function as cultural emblems that intertwine craftsmanship, heritage and collective pride. Translating gastronomic texts therefore requires attention not only to descriptive accuracy but also to their aesthetic, symbolic and affective dimensions.

In Chinese culture, food similarly operates as a vehicle of philosophical reflection, artistic expression and emotional resonance. Culinary images often serve as metaphors for sensorial, moral and even cosmological experiences (Song 2025). This shared symbolic richness means that gastronomic translation between Spanish and Chinese involves bridging two cultural universes that rarely align perfectly. The translator must mediate between these divergent worldviews, reconstructing meaning in ways that preserve cultural depth and interpretive nuance.

2.2 Culinary lexicon, linguistic tradition and cultural load

The Spanish gastronomic lexicon has developed between everyday use and specialized culinary discourse (Eberenz 2014). Many terms such as *salazón*, *curado*, *perfilado*, encapsulate artisanal processes shaped by centuries of cultural practice. Torres (2014) points out that numerous culinary expressions have entered normative dictionaries, becoming part of the cultural heritage of Spanish-speaking communities.

These culturally anchored terms lack direct equivalents in Chinese. Their translation requires negotiation, contextualization and interpretive choices. As Moreno (2021) argues, etymological layers embed each culinary term in a

network of historical and symbolic associations. AI models, however, do not access cultural memory; they generate statistically plausible outputs that often flatten symbolic meaning. As a result, they may produce lexically correct but culturally impoverished translations.

For this reason, gastronomic translation demands a human interpretive capacity capable of recognizing implicit references, sensory imagery and cultural logic. The translator's role becomes one of cultural preservation, ensuring that the aesthetic and symbolic dimensions of food discourse remain intelligible and resonant for the target audience.

2.3 Machine translation, AI and cultural mediation

Recent advances in neural and generative models have significantly improved machine translation quality, especially in syntax and contextual cohesion. Automatic systems offer increasingly coherent renderings, yet they continue to face limitations with culturally saturated content. As Bowker (2020) observes, these models often reproduce surface-level textual patterns while missing culturally encoded meanings, implicit references and rhetorical intent. The result is a neutralized tone and a loss of cultural resonance.

Li (2024) conceptualizes gastronomic translation as an act of intercultural mediation that balances cultural fidelity with communicative clarity. This requires accuracy, integrity and adaptive strategies that preserve historical and aesthetic value. In this view, the translator operates as a symbolic mediator who interprets food discourse as a carrier of memory and identity.

Similarly, Zhang/Jiao (2025) regards the translation of culinary products as a dynamic process of intercultural adaptation. The translator negotiates meanings, identities and expectations in order to harmonize authenticity and accessibility. Human translators take context into account, whereas AI works within rigid algorithmic boundaries that tend to dilute cultural diversity.

3 Methodology

This study adopts an applied qualitative approach based on the comparison between machine-generated translations and human translations produced in a university classroom context. The aim is to observe how AI-driven tools reproduce, or neutralize, cultural, symbolic and discursive elements in gastronomic texts, and how human mediation restores these dimensions in the target language.

The methodological framework follows Li's (2024) and Zhang/Jiao's (2025) understanding of translation as an act of intercultural mediation. From this perspective, translators do not merely transfer linguistic meaning; they reconstruct cultural experience and communicative intent. This paradigm allows us to analyze how human translators interpret gastronomic discourse and how their strategies differ from algorithmic decision-making.

3.1 Corpus and text selection

The analysis is based on a corpus composed of Spanish descriptions of gourmet products, with a central focus on Iberian ham due to its strong cultural symbolism and relevance in tourism promotion. These texts were translated into Chinese using five neural and generative systems: DeepL, ChatGPT, Kimi, DeepSeek and Copilot; each representing different technological and cultural ecosystems.

For comparison purposes, a reference human translation was included in the corpus. This translation was produced by undergraduate students enrolled in a Spanish–Chinese translation course during the second semester of the 2024/2025 academic year.

This dual corpus, machine-generated and human-produced, enabled a layered comparison that highlights differing conceptions of translation, interpretation, and cultural representation.

3.2 Machine translation and comparative criteria

To analyze machine translation, five neural and generative language models were employed: DeepL, ChatGPT, Kimi, DeepSeek, and Copilot. These systems originate from distinct technological and cultural environments. Each source text segment was translated into Mandarin Chinese, a process that yielded multiple automated versions for subsequent comparison against a human reference translation.

The comparison was conducted based on five principal criteria: terminological treatment, appropriateness of commercial tone, translation of metaphors and cultural references, linguistic fluency and naturalness, and discursive coherence in the target language.

This study adopts a comparative approach, contrasting the solutions proposed by AI models with versions produced by university students. This contrast reveals differences between two modes of understanding and resolving translation challenges: one relies on algorithmic operations driven by contextual predictions, while the other is grounded in cultural interpretation and the translator's discursive sensitivity.

From this perspective, the objective is not to measure the technical efficacy of Artificial Intelligence, but to understand how human reasoning reconstructs the cultural meaning that automation often flattens. This analysis highlights the indispensable value of human mediation in translating gastronomic discourse.

3.3 Classroom training experience

The practice was conducted in the course “*Translation for Humanities, Social Sciences, and Legal Studies A–C (Spanish-Chinese)*” during the second semester of the 2024/2025 academic year. Students completed a direct translation into Chinese of the text “*El jamón ibérico, un producto gourmet*” (*Iberian Ham, a Gourmet Product*) and subsequently prepared a reflective report explaining the strategies and techniques they had employed.

Qualitative analysis of the translations and reports revealed the use of cultural domestication strategies, semantic amplification, and stylistic

recreation, while also assessing the participants' developed level of intercultural awareness. This activity became a small translation laboratory where human competence and AI-generated solutions could be directly compared.

Subsequently, the students' translation results were compared with the versions produced by the AI models. This aimed not only to highlight differences in methodology but also to identify the underlying translation philosophies activated by each approach. Thus, this classroom practice served as a space for reflection where human intelligence and AI engaged in a dialogue concerning culture, language, and translational creativity.

3.4 Study evaluation and scope

Although this study relies on a limited sample and is exploratory in nature, its value lies in the integration of three distinct perspectives. First, the comparative analysis of translated texts. Second, the observation of the formative process within the classroom. Finally, critical reflection on the importance of AI analysis within cultural mediation. The convergence of these three perspectives enables the identification of cultural gaps present in machine translation, while also demonstrating how the university classroom can become a space for reflection where AI serves as a pedagogical tool to cultivate critical, creative, and culturally competent translators.

This dialogue between human and automated translations forms the core of the evaluation, illustrating how the translator's cultural mediation compensates for the interpretive limitations of AI. Simultaneously, it opens new avenues for reflecting on technology's evolving role in both translator education and professional practice.

4 Comparative analysis and classroom results

4.1 General approach to the analysis

The comparative analysis between machine and human translations examines how cultural, sensory, and symbolic elements in gastronomic texts are interpreted and transferred. Rather than focusing solely on linguistic accuracy or formal equivalence, this study seeks to understand the translation decisions that shape the transmission of cultural meaning and the expressive tone of the original text.

Gastronomic translation requires competencies extending beyond purely linguistic mastery. We observe how errors in AI output and solutions provided by students reveal varying degrees of cultural mediation. This analysis is not limited to identifying errors or successes, but aims to pinpoint the cognitive differences distinguishing these two translation modes. Where AI processes text through context predictions based on large corpora, human translators employ more interpretive reasoning, cultural empathy, and more conscious communicative intent.

Comparing these two approaches will enable us to understand the patterns of neutralization characteristic of machine translation, as well as the cultural adaptation strategies employed by students in the following sections.

4.2 Machine translation: neutralization and literalism

The analysis of outputs from various AI models reveals that while these tools generally produce coherent and accurate translations through contextual prediction and neural learning mechanisms, their algorithmic approach, divorced from genuine cultural experience, tends toward excessive literalism, stripping symbols of their deeper meaning. Rico (2017) emphasizes that machine translation should be understood as a cyclical process where the translator manages the technology rather than merely operating it. From an

educational perspective, this requires placing students at the centre of learning and fostering critical and collaborative competencies that integrate AI as a support for their own capabilities.

The analyzed tools also share a tendency to prioritize structural coherence and lexical precision over the symbolic or affective weight of the text. Although the resulting translations are formally correct, their reading reveals a neutralization of the promotional tone and a significant loss of cultural nuance.

For example, the text's opening sentence:

“El jamón ibérico nos hace más felices, ya lo dijo Camilo José Cela: ‘Es un bocado propio de bienaventurados.’”

was translated literally by several automated systems as:

“伊比利亚火腿让我们更幸福。卡米洛·何塞·塞拉早已说过：‘这是属于有福之人的一口食物’。”

Despite its formal correctness, this version lacks the poetic resonance and emotional tone of the original. The expression “一口食物” (*“a mouthful of food”*) appears mechanical and semantically flat, while “有福之人” (*“fortunate person”*) fails to capture the spiritual connotation of “bienaventurados”. The result is a textually correct but emotionally diminished version.

Similarly, technical terms such as “*perfilado, asentado or curado*” are translated accurately but without the narrative component that connects them to artisanal tradition and the concepts of time and heritage. Overall, machine translations tend to construct an informative discourse, more descriptive than evocative, that fails to reproduce the sensory tone or cultural identity of the source text. This loss of symbolic density contrasts sharply, as will be seen in the following section, with the adaptation and recreation strategies applied by human translators.

4.3 Human translation in the classroom: strategies and cultural mediation

In contrast to the neutrality characteristic of algorithmic processing, the translations produced by students demonstrate greater cultural and stylistic adaptation. These versions clearly highlight the translator’s capacity for mediation – essential for grasping the profound communicative intent of the source text and reconstructing its aesthetic value within the target culture. The crucial distinction lies in the fact that, unlike the AI’s predictive model based on statistical patterns, human translation operates through contextual interpretation, cultural inference, and conscious discursive adaptation.

As shown in the examples below, the students went beyond simple lexical substitution and actively reconstructed the cultural experience associated with the product.

Table 1: Machine Translation vs. Human Cultural Interpretation. (Author’s own work)

Original Fragment	Machine Translation	Students’ Translation	Type of Intelligence	Observation
“El jamón ibérico nos hace más felices, ya lo dijo Camilo José Cela: ‘Es un bocado propio de bienaventurados.’”	“伊比利亚火腿让我们更幸福。卡米洛·何塞·塞拉早已说过：‘这是属于有福之人的一口食物’。”	“伊比利亚火腿能为人们带来无上欢愉，正如诺贝尔文学奖得主卡米洛·何塞·塞拉曾赞叹道：‘此味只应天上有’。”	Algorithmic Processing: statistical prediction devoid of cultural contextualization. Interpretive Processing: symbolic and aesthetic reconstruction.	While AI output retains formal correctness, the students’ translations converted the text into an experience rich in cultural and emotional resonance.

In this example, we observe that students employ strategies of idiomatic recreation, such as using localized set phrases and proverbs like “无上欢愉、此味只应天上有” (supreme pleasure, a flavor that should only exist in heaven),

or explanatory amplifications like “正如诺贝尔文学奖得主...” (as noted by the Nobel Prize in Literature winner ...), to preserve the text’s emotional intensity and adapt it to the poetic sensibilities of the Chinese reader. Far from compromising fidelity to the original, this approach not only enhances the translational value but also ensures full communicative effectiveness.

Another significant case emerges in the expression “forma parte de una herencia” (“forms part of a heritage”). Most AI systems produced literal versions like “属于一种遗产” (“belongs to a type of heritage”), whereas several students rephrased it as “升华为传承千年的文化瑰宝” (“sublimes into a cultural treasure passed down through millennia”). This semantic amplification not only preserves the notion of legacy but integrates concepts of time, tradition, and national pride, which resonate deeply within the target culture.

Similarly, in technical sections of the text, such as *perfilado* (trimming), *salazón* (salting), *curado* (curing), the students’ translations demonstrate mastery of specialized terminology enriched with expressive nuances absent in AI outputs. Instead of “修整形状” (“trimming the shape”), some students chose “匠人运刀如笔” (“the artisan wields the knife as a calligrapher wields a brush”), thereby emphasizing the aesthetic and artisanal dimensions of the process.

Collectively, the students’ translations reveal a clear intercultural sensitivity. They transcend mere word-for-word transfer by actively interpreting, negotiating, and recreating meanings with consideration for the target reader and the product’s symbolic value.

The strategies they apply, such as domestication, amplification, and idiomatic recreation, align with the typologies described by Guan (2023) in his study on translating Chinese dishes. In that work, the author identifies four primary procedures: literal translation, free translation, phonetic translation, and regional translation. Although the cultural contexts differ, the objective remains identical: to preserve both the intelligibility and cultural authenticity of the gourmet product. This convergence suggests that human strategies stem from universal cognitive patterns linked to cultural mediation. In contrast, AI tends to replicate a mechanical literality that fails to attain this level of interpretative depth.

4.4 AI–student comparison: implications for translation studies

Translations produced by AI models demonstrate a clear pattern of cultural neutralization. They exhibit grammatical correctness and intrinsic textual coherence, yet maintain a uniform tone that omits many of the symbolic and cultural nuances inherent to gastronomic texts.

Across the thirteen translations produced by students, we observe the application of an interpretive intelligence oriented toward cultural mediation. They do not merely transfer words. They amplify meanings, adapt idiomatic expressions, rephrase metaphors, and preserve the aesthetic and emotional dimension of the source text.

A clear distinction emerges between these two approaches to translation. AI processes text as structured data, whereas the human translator approaches it as a process requiring interpretation and contextualization. Machine translators achieve notable formal precision but lack adequate cultural understanding. The human translator, however, employs a discursive sensitivity that enables them to act as a symbolic mediator between different cultures.

This difference manifests across several levels. Lexically, automated models reproduce terminology accurately but fail to convey its affective nuances. The students, in contrast, rephrase and amplify with an aesthetic sensibility beyond the system's reach. Regarding tone, AI often strips the text of its promotional subjectivity, while students preserve this dimension through rhetorical and figurative devices. Structurally, the models tend to maintain a mechanical linearity, whereas the translators reorganize the discourse to achieve natural fluency in Chinese, thereby generating a persuasive coherence that respects the original's communicative function.

In essence, the value of human translation lies not only in its formal correctness but also in its capacity to transmit cultural and emotional experiences. Hence, the translator remains an indispensable mediator even in this increasingly automated era. Cultural empathy and human creativity remain distinctly human attributes in the act of translation, crucial for preserving authenticity in intercultural communication.

4.5 Pedagogical outcomes and critical reflection

Consequently, we have observed how both human and automated translation have served a clear pedagogical purpose. Students gained first-hand experience in recognizing the difference between translating with a focus on cultural dimensions and translating through algorithmic prediction. This practice strengthened their meta-translation competence and deepened their understanding of the broader implications of their own interpretive choices.

From a teaching perspective, this initiative proved to be a valuable tool for fostering deep reflection on translation and the translator's role as a cultural bridge. It prompted a clear awareness of the boundaries of automation and highlighted the irreplaceable value of the human element in conveying messages effectively into the target language.

Throughout this exercise, students identified the limitations of AI in handling cultural references and idiomatic expressions. This discovery helped them recognize the need for their own interpretive and ethical intervention. They thus understood that AI is not a substitute for the translator, but rather an instrument for analysis and self-assessment that helps refine linguistic, discursive, and symbolic judgment.

Furthermore, the activity facilitated the identification of key professional competencies for translators: intercultural adaptation ability, stylistic and discursive awareness, and the ethical responsibility in managing cultural meaning.

In conclusion, the objective was not to pit machine against human, but to place these automated models at the service of student learning. In this way, future professionals will be equipped to engage with technology critically and reflectively, without relinquishing their primary role as cultural mediators.

5 Discussion

Based on these findings, we examine the translational, cultural, and educational implications arising from the comparison between machine-generated versions and those produced by students. While AI tools achieve formal accuracy, they

demonstrate a clear tendency to neutralize or suppress symbolic elements and cultural connotations. In contrast, the students' versions exhibit greater creativity, more refined aesthetic sensitivity, and a superior ability to reconstruct the communicative and emotional intent of the gastronomic descriptions.

These results confirm that the human translator does not merely transfer meaning; rather, they recreate cultural experiences, a process that transcends semantic equivalence and remains beyond the current capabilities of AI. These findings align with Wang/Gao (2025), Yang/Qin (2024), and Huang/Xu (2024), who emphasize that gastronomic translation demands symbolic and experiential recreation connected to embodiment, gustatory memory, and cultural sensitivity. The translator's mediation, grounded in sensory and symbolic dimensions, preserves the cultural substance of the text and ensures its communicative efficacy.

This pedagogical practice, developed in the university classroom, demonstrates how the critical and reflective use of AI can serve as a significant educational resource. By comparing both versions, students can understand the limitations of technology, appreciate their own interpretive capacity, and develop meta-translational competence that integrates analysis, creativity, and ethics. The classroom thus becomes a transformative space where experimentation with new technological tools serves as a mechanism for active learning and a stimulus for cultural reflection.

In this regard, Li (2025) highlights the need to orient translation studies toward collaborative models between technology and the humanities, utilizing specialized corpora and interdisciplinary methodologies. University education must equip professionals to integrate digital tools without sacrificing the humanistic core that gives meaning to their profession.

For her part, Zhang/Jiao (2025) conceptualizes gastronomic translation as a dynamic process of intercultural adaptation in which the translator constantly negotiates between cultural authenticity and reader comprehension. Translation involves modulating values, symbols, and communicative practices to preserve the vitality of meaning across diverse cultural contexts. This perspective confirms that translation, far from being a mechanical operation, is an act of living, situated mediation where human sensitivity remains indispensable.

Therefore, these results make it evident that AI cannot replace the translator's understanding but rather complements it from a perspective that fosters reflection and creativity. It is the mission of the contemporary translator to maintain balance between technological efficiency and cultural sensitivity, recognizing that their task involves not only linguistic transfer but also serving as a symbolic mediator with intercultural responsibility. With this perspective, AI can undoubtedly be integrated into both education and professional practice without compromising the humanistic dimension of the translation act.

6 Conclusions

Throughout this chapter, we have shown that current AI-driven translation technologies have achieved notable progress in managing structural, terminological and contextual information. Yet, despite these advances, they continue to lack the cultural and aesthetic sensitivity required to reproduce the symbolic, sensory and experiential dimensions intrinsic to gastronomic discourse. By contrast, human translators demonstrate interpretive, adaptive and creative strategies that successfully preserve the cultural identity, emotional resonance and promotional intent of the source text. Together, these findings underscore that translational competence extends far beyond linguistic accuracy: it requires contextual awareness, intercultural understanding and the ability to reconstruct meaning in culturally meaningful ways.

The classroom-based practice designed to promote a reflective and critical use of AI proved to be a highly valuable educational tool. By confronting machine-generated outputs with their own translations, students were able to recognize both the potential and the limitations of automated systems, while also deepening their understanding of translation as an interpretive, culturally situated activity. As a result, the classroom becomes an experimental environment in which human and Artificial Intelligences engage in dialogue, transforming AI from a perceived threat into a pedagogical ally that fosters autonomy, creativity and ethical awareness.

While AI produces structurally coherent renditions, it systematically reduces discourse to surface-level meaning, stripping texts of their symbolic, sensory and affective layers. Faced with this limitation, students' translations consistently revealed an ability to reconstruct cultural references through strategies such as semantic amplification, metaphorical adaptation and stylistic recreation. The fundamental divergence between both types of output does not arise from grammatical errors, but from the loss of experiential depth and cultural fidelity, elements that define the persuasive and identity-laden nature of gastronomic discourse.

The evidence gathered therefore affirms a crucial distinction between algorithmic precision and cultural mediation. The latter remains the hallmark of human translation and a requirement that current AI systems cannot yet emulate. Ultimately, it is the human capacity to interpret, negotiate and culturally re-embody meaning that ensures the communicative authenticity of the target text, particularly in domains where meaning is inseparable from tradition, identity and sensory experience.

In conclusion, while AI can support translation practice and enrich translation pedagogy, it cannot replace the human translator's essential role as cultural mediator. In fields such as gastronomy and tourism, where discourse is inevitably tied to heritage, emotion and symbolic value, the depth, authenticity and cultural integrity of the translation depend fundamentally on human interpretive intelligence.

References

- BOWKER, L. (2020): *Machine Translation and Global Research: Towards Improved Machine Translation Literacy in the Scholarly Community*. Bingley, UK: Emerald Publishing.
- CARRETERO DOMÍNGUEZ, M. Á. (2014): *Antropología culinaria* [Culinary anthropology]. Puebla, México: Universidad Popular Autónoma del Estado de Puebla.
- EBERENZ, R. (2014): "El léxico español de la alimentación y la culinaria en su historia: Fuentes y líneas de investigación" [The Spanish Lexicon of Food and Culinary Practices in its History: Sources and Lines of Research]. *Dándole cuerda al reloj*

- [*Winding the Clock*]. Eds. V. Álvarez / E. Díez del Corral / N. Reynaud (23–46). Valencia, Spain: Tirant lo Blanch.
- GUAN, Q. (2023): “中国特色菜肴的英文翻译策略研究” [Strategies for Translating Chinese Specialty Dishes into English]. *Food Frontier Research*: 82–84. Beijing, China.
- HUANG, Q. & XU, J. (2024): “中西饮食文化差异在英译中的体现” [Manifestations of Cultural Differences between Chinese and Western Cuisines in English Translation]. *Food and Machinery* 40.2: 842–845. Beijing, China.
- LI, G. (2025): “中国饮食文化翻译研究的回顾与前瞻” [A Review and Prospects of Research on the Translation of Chinese Food Culture]. *Journal of Luoyang Normal University* 44.1: 93–97. Luoyang, China. DOI: <https://doi.org/10.16594/j.cnki.41-1302/g4.2025.01.001>.
- LI, J. (2024): “跨文化视角下的饮食文化英译研究” [A Study on the English Translation of Food Culture from an Intercultural Perspective]. *China Food* 39.2: 138–140. Beijing, China.
- LÓPEZ GARCÍA, J., JUÁREZ, L. M. & MEDINA, F. X. (2016): “Usos y significados contemporáneos de la comida desde la antropología de la alimentación en América Latina y España.” [Contemporary Uses and Meanings of Food from an Anthropological Perspective in Latin America and Spain]. *Revista de Dialectología y Tradiciones Populares* 71.2: 327–370. Madrid, Spain. DOI: <https://doi.org/10.3989/rdtp.2016.02.007>
- LU, C. (2023): “饮食道途：英文小说《品味中国》对中国饮食文化的翻译再现与海外接受” [The Path of Gastronomy: English Translation and Overseas Reception of the Novel *Tasting China*]. *Foreign Literature Studies* 45.4: 102–110. Wuhan, China.
- MORENO MORENO, M. Á. (2021): “La microestructura del diccionario etimológico: Los primeros pasos de esta práctica lexicográfica.” [The Microstructure of the Etymological Dictionary: Early Steps in this Lexicographical Practice]. *Revista de Lexicografía* 27: 259–270. A Coruña, Spain.
- RICO, C. (2017): “La formació de traductores en traducció automàtica.” [Training Translators in Machine Translation]. *Tradumàtica: Tecnologies de la Traducció* 15: 75–96. Barcelona, Spain. DOI: <https://doi.org/10.5565/rev/tradumatica.200>
- SONG, Q. (2025): “古诗词中饮食文化意象的英文翻译策略研究” [Strategies for Translating Food-Related Cultural Imagery in Classical Chinese Poetry into English]. *Food and Machinery* 41.5: 249–252. Beijing, China.

- TORRES MARTÍNEZ, M. (2014): “Léxico culinario autorizado en el primer diccionario académico.” [Culinary Lexicon Included in the First Academic Dictionary]. *Cuadernos del Instituto de Historia de la Lengua* 9: 295–324. San Millán de la Cogolla, Spain.
- WANG, X. & GAO, Y. (2025): “体认翻译学视角下的山西美食翻译研究” [Translation of Shanxi Cuisine from the Perspective of Embodied Translation Studies]. *Masterpieces Review* 27: 30–33. Taiyuan, China. DOI: <https://doi.org/10.20273/j.cnki.mzxs.2025.27.031>
- YANG, C. & QIN, H. (2024): “基于体认翻译学的中医药膳翻译研究——以四川菜博物馆为例” [Translation of Traditional Medicinal Cuisine from the Perspective of Embodied Translation Studies: A Case Study of the Sichuan Cuisine Museum]. *Overseas English* 2024.2: 37–39. Harbin, China.
- ZHANG, L. & JIAO, J. (2025): “跨文化交际背景下中餐小吃菜名英译研究——以《美食译苑——中文菜单英文译法》为例” [English Translation of Chinese Snack Dish Names in Intercultural Communication: A Case Study of Meishi Yiyuan]. *Xin Chuanqi (新传奇)* 2025(16): 41–43.

Textual Molds for In-Context Learning at Sea

Can Generative AI Produce Shipping Forecasts without Semantic Guidance?

Abstract

Seewetterberichte stellen ein stark eingeschränktes Genre der Fachkommunikation dar. Sie sind auf operative Effizienz und Sicherheit ausgelegt und unterliegen strengen Normen. Sie zeichnen sich durch eine kontrollierte Sprache, feste lexikalisch-grammatikalische Strukturen und semantische Vorhersehbarkeit aus und weisen das auf, was Sinclair (1991) als Textmuster und Condamines (2018) als Diskursfixiertheit bezeichnen. Der Beitrag untersucht, ob LLMs diese strukturellen und konzeptuellen Einschränkungen replizieren können. Es wurde zuerst eine Korpusanalyse französischer Vorhersagen durchgeführt, um wiederkehrende lexikalische und syntaktische Muster zu identifizieren, die das Genre definieren. Sie wurden dann mit den zugrunde liegenden semantischen Strukturen artikuliert. Diese Muster dienen schließlich der Bewertung der Ergebnisse von zwei LLMs (Claude-opus-4 und Grok-4-fast), die unter Zero-Shot-Bedingungen generiert wurden. Die Ergebnisse zeigen, dass beide Modelle zwar oberflächliche Regelmäßigkeiten erfassen, jedoch häufig von den domänenspezifischen

.....
1 SARAH THEROINE Centre Interlangues Texte Image Langage (UR 4182), Université Bourgogne Europe; LAURENT GAUTIER Centre Interlangues Texte Image Langage (UR 4182), Université Bourgogne Europe; CHRISTOPHE CRUZ Laboratoire Interdisciplinaire Carnot de Bourgogne – UMR CNRS 6303, Université Bourgogne Europe; GENNARO ASCIONE Centre Interlangues Texte Image Langage (UR 4182), Université Bourgogne Europe, Università degli Studi di Napoli Parthenope

konzeptuellen Strukturen abweichen. Dies wirft weitergehende Fragen hinsichtlich der Genre-Sensitivität, der epistemischen Kohärenz und der Grenzen der promptbasierten Textgenerierung für spezialisierte Kommunikationsaufgaben auf.

Keywords AI, linguistic patterns, LSP, shipping forecast, textual mold

1 Introduction

Shipping forecasts represent one of the most highly codified genres within professional communication, characterized by standardized phraseology, condensed syntax, and predictable surface structures. As a type of controlled language, their narrow lexical range and grammar ensure clarity, safety, and international intelligibility in high-stakes contexts (Bocanegra-Valle 2011, 2015; Condamines 2018). This study focuses on French shipping forecasts issued by Météo-France. These can be considered a subtype of specialized discourse shaped by institutional and international regulations, including the IMO² and WMO³ standards governing the Global Maritime Distress and Safety System (GMDSS)⁴.

This study is guided by the following research question: To what extent can LLMs reproduce both the surface patterns and the underlying conceptual constraints of shipping forecasts? Due to their formulaic nature and semantic regularity (Gautier 2008), these texts offer a robust case study to evaluate the genre sensitivity of Large Language Models (LLMs). This paper explores whether two state-of-the-art models — Claude-opus-4 (dense architecture) and Grok-4-fast (Mixture-of-Experts) — can replicate both the surface and conceptual structures of shipping forecasts in zero-shot conditions. Our method therefore combines corpus linguistics, frame semantics and computational

.....
2 International Maritime Organization (IMO): <https://www.imo.org/>

3 World Meteorological Organization: <https://wmo.int/about-us/governance/technical-commissions/standards-and-recommended-practices>

4 Global Maritime Distress and Safety System: <https://gmdss.com/>

text generation. A manually annotated corpus of Météo-France forecasts was used to identify recurrent patterns and core frame elements (FEs). The LLMs were then prompted to generate French forecasts, allowing us to assess their ability to reproduce constrained, domain-specific discourse without explicit semantic input.

This article is structured as follows: Section 2 briefly reviews relevant literature, section 3 outlines the methodology applied, section 4 presents the experimental setup, section 5 reports findings, and section 6 discusses implications and future directions.

2 State-of-the-art

Shipping forecasts are structured sequences of fixed lexico-grammatical patterns, each serving distinct semantic and discursive functions. Bocanegra-Valle (2015) describes them as “straightforward and highly conventionalized texts”, thus that such texts are patterned, standardized and highly conventional, a perspective that can be further refined by emphasizing their sequential organization. Rooted in operational routines that prioritize clarity and safety, these texts exemplify controlled language, characterized by restricted vocabulary, elliptical syntax, and minimal variation (Condamines 2018; Johnson 1995) all the while benefitting from scientific linguistic knowhow. The concept of usability, central in ergonomics as it puts people at the heart of needs, is adapted to the perspective of ergonomic linguistics. In order to put this branch of linguistics to work, we need to take into consideration on the one hand the real life corpora or language use in context, and on the other hand implement tests to evaluate their proper usage and relevance. We base our study on controlled languages because of the way they are constituted (the rules that make them). Despite their constraints, linguistic research has rarely addressed their patterned structure as organized templates. Sinclair’s textual mold (1991) offers a productive lens for analyzing this genre, framing forecasts as epistemically structured templates that stabilize production and interpretation (Condamines 2018). Recent studies formalize these molds

through pattern extraction and semantic modeling, showing their link to domain-specific knowledge structures (Aussenac-Gilles/Condamines 2012; Orna-Montesinos 2012). As textual molds intertwine rigid phraseological patterns with domain-specific conceptual constraints, they present a particularly revealing challenge for generative models, which may imitate surface regularities without internalizing the underlying semantic architecture of the mold.

Recent advances in LLMs have renewed interest in replicating such constrained genres. Through in-context learning (ICL), models infer task structure from prompts alone, even in zero-shot conditions, mimicking analogy-based reasoning (Dong et al. 2024). This mechanism is suited to controlled language, allowing models to generalize from minimal structured input. Claude-Opus-4, a dense model optimized for long-context reasoning (200k tokens), emphasizes coherence and structured task performance. In contrast, Grok-4-fast adopts a Mixture-of-Experts (MoE) architecture, distributing input across specialized subnetworks to support faster and more flexible generation (256k tokens). Said models embody complementary approaches to constrained text generation: stable reasoning versus expert-routed contextual specialization.

3 Method

The first stage involved the manual extraction of recurring lexico-grammatical patterns from a corpus of French shipping forecasts. Its aim was to collect stable associations of lexical and syntactic units that characterize the genre. Special attention was devoted to morpho-syntactic constructions linking meteorological phenomena, measurements, movements, and spatio-temporal references, drawing on a constructionist perspective as proposed by Liégeois/Gautier (2025a, 2025b) and Liégeois et al. (2023). Said patterns reflect cognitively efficient strategies for encoding specialized knowledge in operational discourse and give major importance to semantic structures. Each pattern was segmented into conceptual elements, allowing for the abstraction of a conceptual sequence. Details on frame-semantic annotation procedures are

not presented here, as the method builds on Theroine et al. (2025). Although not used here for training, these references served to evaluate on an objective basis whether the outputs of LLMs were in line with the genre's formal and semantic regularities. The second stage employed ICL techniques to automate the generation of shipping forecast sentences. In a zero-shot setting, the models were prompted to generate new sentences, relying on the following instructions. This setting has been chosen because it forces the models to infer the genre's structural and semantic constraints from a single authentic instance, making it an effective test of true zero-shot generalization.

“follow the example given to generate 10 new sentences on marine weather \nText:” Dépressions 1007 hPa au large d'Ouessant se creuse et se décale vers le sud-est prévue 1001 hPa sur les pays de la Loire ce soir. Elle se comble sur le centre de la France demain dimanche.
+ “\ngenerated sentence:”

4 Data description

The corpus consists of authentic shipping forecasts published by Météo-France. As specified in the introduction, such forecasts represent a recurrent and highly standardized form of expert communication, designed to ensure navigational safety and support decision-making at sea. Issued regularly for specific maritime zones, they follow an established discursive routine that promotes stability, precision, and interpretive predictability. Pritchard/Kalogjera (2000: 185) observe that these texts are characterized by “a limited vocabulary and simple grammatical structure to suit the specific requirements”. The latter being: safety, reliability, and uniformity. Each forecast adheres to a fixed macrostructural sequence composed of successive informational sections (cf. Figure 1), or moves, corresponding to the communicative stages (Bocanegra-Valle 2015).

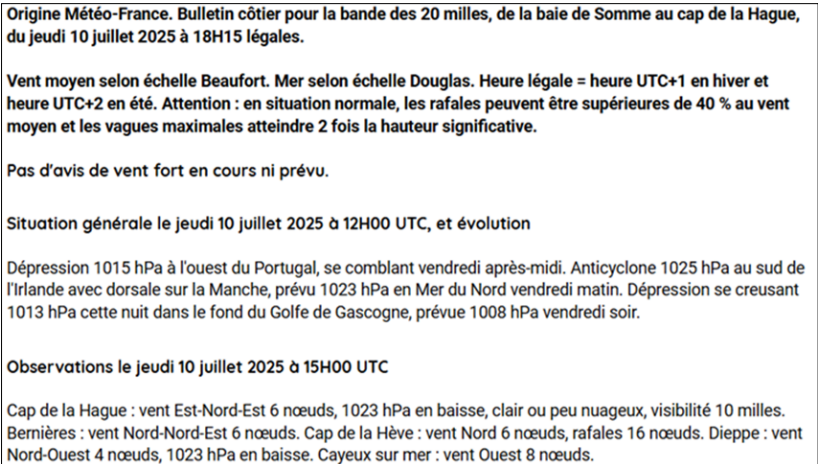


Figure 1: Macrostructure of marine weather forecast

The corpus compiled consists of written shipping forecasts collected between February and March 2024, covering the Western Mediterranean sea area.

Number of documents	Number of tokens	Number of types	Number of sentences
29	38,416	550	3,245

Figure 2: Quantitative information about the corpus

5 Results

5.1 Identifying standardized patterns and semantic frames

Shipping forecasts follow a constrained style, marked by elliptical syntax, nominal condensation, and frequent lexico-grammatical repetition. Verbs are rare (10 %) and mostly appear as participles (Theroine et al. 2025). Meaning is conveyed through dense noun and preposition phrases (cf. Figure 3) encoding space, time, and phenomena. This structure reflects the principles of controlled

languages (Condamines 2018) and Sinclair’s (1991) notion of a textual mold shaped by institutional constraints.

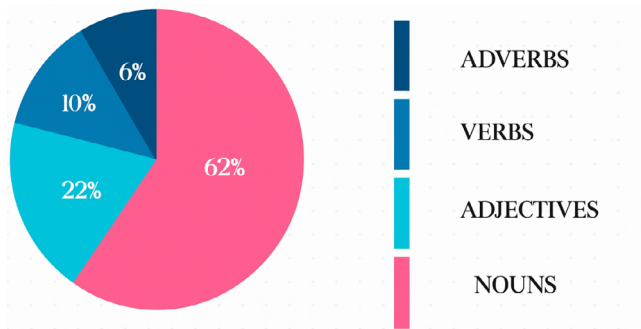


Figure 3: Distribution of grammatical classes in the corpus

Predicative verbs and their combinatory profiles (Gautier 2014) were extracted using Sketch Engine (Kilgariff et al. 2004). Through targeted Corpus Query Language (CQL) searches, predicates were identified and mapped onto semantic frames following the FrameNet typology (Fillmore 2006). A methodological decision was to center CQL queries around the fixed expression *prévue* (“forecasted” or “expected”). This lexical item, although morphologically derived from *prévoir* (“to forecast” or “to anticipate”), has undergone semantic bleaching and grammaticalization: Rather than functioning as an inflected verb, *prévue* acts as a disciplinary figement, a stable discourse marker, signaling the transition from ongoing dynamics (*se creuser* – “to deepen”, *se déplacer* – “to move”) toward expected atmospheric states (*prévue 1001 hPa ce soir* – “forecasted 1001 hPa this evening”).

French sentences	Translated versions
<i>Dépresseions 1007 hPa au large d'Ouessant se creuse et se décale vers le sud-est prévue 1001 hPa sur les pays de la Loire ce soir. Elle se comble sur le centre de la France demain dimanche</i>	Low-pressure system at 1007 hPa off the coast of Ouessant is deepening and shifting southeastward, forecasted at 1001 hPa over the Pays de la Loire this evening. It will fill over central France tomorrow, Sunday.
<i>Dépresseions à 1007 hPa près des côtes bretonnes se creusant et se déplaçant vers le sud-est dans les prochaines heures, prévue 1001 hPa sur la région des Pays de la Loire ce soir. Elles s'affaiblissent progressivement sur la région lundi matin.</i>	Low-pressure systems at 1007 hPa near the Brittany coast deepening and moving southeast over the next few hours, forecasted at 1001 hPa over the Pays de la Loire region this evening. They will gradually weaken over the region Monday morning
<i>Tempêtes à 983 hPa au large de la Manche ce matin se déplaçant vers l'est dans les prochaines heures, prévue à 980 hPa sur la côte normande ce soir. Elles perdent en intensité sur le nord de la France demain matin.</i>	Storms at 983 hPa off the English Channel this morning moving eastward over the next few hours, forecasted at 980 hPa over the Normandy coast this evening. They will weaken over northern France tomorrow morning.

Figure 4: Original weather sentences and their translated versions

In these sentences, *prévue* consistently appears in the same syntactic position and with the same semantic function, confirming its treatment as a non-generic but structurally essential slot, corresponding to the Expectation FE. It serves as a hinge between the description of the current atmospheric evolution and its projected state, thus encoding the prospective dimension crucial to forecast communication (Condamines 2018; L’Homme 2020). Due to its role, *prévue* serves as an effective anchor for extracting low-variability, formulaic sentences. Figure 4 consistently displays six core semantic slots: Phenomenon, Measure, Location, Movement, Time_period, and Expectation.

<phenomenon> <measure> <location> <time_period> <move>
<expectation>

Therefore, based on the identified slots, we manually extracted the following lexico-grammatical patterns (cf. Figure 5).

First annotated sentence	Second annotated sentence	Third annotated sentence
Dépressions [phenomenon] 1007 [measure_value] hPa [measure_unit] au large d'Ouessant [location] se creuse [V. intensity] et se décale [V. Mouvement] vers le sud-est [Location] prévue [expectation] 1001 [measure_value] hPa [measure_unit] sur les pays de la Loire [Location] ce soir [time_period]. Elles [phenomenon] se comblent [V. movement] sur le centre de la France [location] demain dimanche [time_period].	Dépressions [Phenomenon] à 1007 [measure_value] hPa [measure_unit] près des côtes bretonnes [location] se creusant [V. intensity] et se déplaçant [V. Mouvement] vers le sud-est [location] dans les prochaines heures [time_period], prévue [expectation] 1001 hPa [measure] sur la région des Pays de la Loire [location] ce soir [time_period]. Elles [phenomenon] s'affaiblissent [V. intensity] progressivement sur la région [location] lundi matin [time_period].	Tempêtes [phenomenon] à 983 [measure_value] hPa [measure_unit] au large de la Manche [location] ce matin [time_period] se déplaçant [V. movement] vers l'est [location] dans les prochaines heures [time_period], prévue [expectation] à 980 hPa [measure] sur la côte normande [location] ce soir [time_period]. Elles [phenomenon] perdent en intensité [V. intensity] sur le nord de la France [location] demain matin [time_period].

Figure 5: Semantically annotated sentences

From these observations, a generalized semantic structure was derived:

<phenomenon> <measure value> <measure unit> <location> <time_
period> <move> <location> <time_period> <expectation> <measure>
<location> <time_period>.

which is followed by a secondary structure describing the subsequent evolution:

<phenomenon> <move> <location> <time_period>.

The structure does not merely describe a surface form, but rather captures a fully developed cognitive frame, rooted in the way forecasts mentally organize meteorological phenomena. Although the frame structure is generally stable, slight variation exists such as optional time markers and shifts between movement verbs (*se déplacer* – “to move”) and intensity verbs (*se creuser* – “to deepen”, *s'affaiblir* – “to weaken”), with intensity adding specificity to movement. Cognitively, such specifications are more relevant for sailors. Each lexical unit was aligned with its FEs, producing a pattern reflecting both linguistic regularity and conceptual coherence. This dataset served as a controlled-language reference to assess whether, under zero-shot conditions,

LLMs could replicate the structured, semantically constrained discourse under study, despite their limitations in generating epistemically grounded content (Bender et al. 2021).

5.2 Artificial shipping sentences

To assess how LLMs internalize the stylistic and structural regularities of shipping forecasts, both models were prompted to generate synthetic forecast sentences. The goal was to evaluate whether each model could autonomously reproduce the controlled-language features of the genre including precise scalar values, elliptical syntax, and event-driven sequencing.

Sentence No.	Claude-opus-4	Grok-4-fast
1	Dépression 1009 hPa au nord de la Bretagne se renforce et se déplace vers l'est prévue 1003 hPa sur la Manche centrale demain matin. Elle s'évacue vers la Mer du Nord lundi après-midi.	Tempête « Alex » 995 hPa au large de Brest se déplace vers l'est, prévue 990 hPa près des côtes vendéennes cette nuit. Elle s'affaiblit sur l'océan Atlantique demain lundi.
2	Anticyclone 1025 hPa sur l'Atlantique s'étend vers le golfe Gascogne prévu 1022 hPa près des côtes landaises ce soir. Il se maintient sur le sud-ouest mardi.	Front froid 1012 hPa au nord de la Biscaye s'intensifie et progresse vers le sud, atteignant 1008 hPa au large de Bilbao au lever du jour. Il se dissipe sur l'Espagne intérieure mardi après-midi.
3	Dépression 1004 hPa en Méditerranée occidentale se creuse rapidement prévue 999 hPa au sud de la Corse demain soir. Elle remonte vers le Golfe de Gènes mercredi matin.	Dépression méditerranéenne 1003 hPa près de Gênes se creuse rapidement vers l'ouest, prévue 999 hPa sur la côte ligure ce soir. Elle se comble au-dessus des Alpes demain matin.
4	Thalweg 1010 hPa sur les Baléares s'approfondit et se décale vers le nord-est prévu 1006 hPa près de la Sardaigne cette nuit. Il se comble progressivement jeudi.	Anticyclone atlantique 1030 hPa au large des Açores se renforce et s'étend vers le nord-est, prévu 1035 hPa sur les îles Britanniques vendredi soir. Il persiste sur l'Europe occidentale samedi.

Sentence No.	Claude-opus-4	Grok-4-fast
5	Dorsale 1018 hPa sur l'ouest de l'Irlande et se renforce et s'oriente vers le sud prévue 1020 hPa en mer Celtique demain midi. Elle s'affaiblit sur le Golfe de Gascogne vendredi.	Rafales de vent jusqu'à 9 Beaufort au large de Gironde, orientées sud-ouest, s'atténuent progressivement vers le nord. Mer agitée devenant forte sur l'estuaire de la Loire cette nuit.
6	Dépression 1002 hPa au large du Portugal se déplace vers le nord-est prévue 998 hPa près du cap Finistère ce soir. Elle se creuse davantage sur le golfe de Gascogne samedi matin.	Houle de nord-ouest 4 mètres au sud de la Sardaigne se propage vers l'est, atteignant 3,5 mètres près de la Sicile demain midi. Elle diminue sur la mer Ionienne en soirée.
7	Front froid 1012 hPa sur la mer d'Iroise progresse vers le sud-est prévu 1008 hPa sur les côtes vendéennes cette nuit. Il s'atténue sur le centre-ouest dimanche.	Coup de vent 7 à 8 Beaufort au large de la Minorque, prévu pour s'orienter nord-ouest avec des rafales à 10. Visibilité réduite sous averses sur Baléares ce soir.
8	Minimum 1005 hPa en mer Ligure se renforce et reste quasi-stationnaire prévu 1001 hPa au sud de Nice demain après-midi. Il se décale vers la Corse lundi soir.	Dépression « Bruno » 1000 hPa sur la mer Tyrrhénienne se déplace vers le sud-est, prévue 997 hPa au large de la Calabre jeudi matin. Elle fusionne avec un front sur la mer Adriatique vendredi.
9	Anticyclone 1028 hPa centré sur les Açores s'étend vers l'Europe prévue 1024 hPa sur le proche Atlantique ce week-end. Il se maintient sur l'ouest mardi prochain.	Mer très forte avec vagues de 5 mètres au nord de la Corse, causée par un vent de nord-est force 8. Conditions s'améliorant progressivement sur la Ligurie demain après-midi.
10	Dépression 1000 hPa au nord de l'Ecosse se déplace rapidement vers le sud-est prévue 995 hPa en mer du Nord demain matin. Elle s'évacue vers la Baltique mercredi.	Anticyclone subtropical 1025 hPa au large de Madère se décale vers le nord, prévu 1028 hPa sur les Canaries ce week-end. Temps stable et vent faible sur l'Atlantique est dimanche.

Figure 6: Results of Claude and Grok

5.2.1 Claude-opus-4: Canonical structure and genre fidelity

When prompted, Claude generates meteorological statements that, at first glance, closely emulate the rhetorical and phraseological conventions of French shipping forecasts. Several sentences display strong conformity to the genre's canonical sequencing: identification of the phenomenon, indication of central

pressure, localization, process predicate, and projected evolution, consistent with the structural regularities identified in corpus-driven studies of the field.

Sentence 1 (*Dépression 1009 hPa au nord de la Bretagne se renforce et se déplace vers l'est [...]* – “Low-pressure system at 1009 hPa north of Brittany is strengthening and moving eastward [...]”) mirrors authentic formulations in which low-pressure systems are tracked through pressure evolution and displacement. The subsequent part (*Elle s'évacue vers la Mer du Nord [...]* – “It evacuates toward the North Sea”) reproduces the standardized expressions found in the corpus. In Sentence 3, Claude produces a genre-faithful representation of a low (*Dépression 1004 hPa en Méditerranée occidentale se creuse rapidement [...]* – “Low-pressure system at 1004 hPa in the western Mediterranean is deepening rapidly”), echoing a pattern observed in the corpus, where Mediterranean or Genoa lows *se creusent* (“deepen”) and *remontent* (“move”) toward the Gulf of Genoa. Likewise, the treatment of *front froid* (“cold fronts”) in Sentence 7 and *minimum dépressionnaire* (“minima lows”) in Sentence 8 demonstrated Claude’s ability to mirror the genre’s specialized terminology and fixed lexico-grammatical patterns.

However, Claude’s apparent mastery of rhetorical templates masks significant conceptual inaccuracies arising from the absence of an underlying meteorological ontology. In Sentence 4 (*Thalweg 1010 hPa [...] s'approfondit et se décale vers le nord-est [...]* – “Trough at 1010 hPa [...] is deepening and shifting northeastward”), Claude assigns a pressure value to a *thalweg* (trough), even though a *thalweg* does not denote a pressure center but rather an elongated axis of low pressure. As shown in the corpus, troughs are described in terms of their orientation and extension, not their central pressure. A similar issue emerges in Sentence 5, where Claude describes a *dorsale* (“ridge”) as *s'orientant vers le sud* (“shifting southward”), a structure that is grammatically plausible yet uncommon in operational forecasts. Typically, ridges are described as *s'étendant* (“extending”) or *s'allongeant* (“lengthening”), not as *orientés* (“oriented”) like mobile entities. These deviations exemplify the type of frame mismatch described in frame semantics (Fillmore/Baker 2012; Fillmore 1976), in which lexical units evoke incompatible sets of FEs. As a result, the model fills in semantic roles (process, measure, movement) that the term’s “true” semantic type does not warrant. Claude introduces non-attested pressure trajectories,

as in Sentence 2 and Sentence 9, where anticyclones (high-pressure systems) are assigned fluctuating pressures that are inconsistent with the slow temporal evolution observed in the corpus. This underscores the observation that the model reproduces surface templates without integrating domain-specific knowledge frames.

Although Claude mimics the rhetorical and phraseological structure of shipping forecasts, it fails to take account of the conceptual distinctions between key meteorological entities. In line with Frame-Based Terminology (L’Homme 2020; Faber 2015), these errors reveal that specialized terms cannot be interchanged without compromising semantic coherence of their argument structure. The model often breaks the mapping between linguistic form and expert knowledge, as observed in specialized genres (Orna-Montesinos 2012). This highlights the limits of template-based generation when semantic constraints are absent.

5.2.2 Grok-4-fast: Semantic fluency and discursive expansion

When prompted, Grok shows a mixed ability to replicate the rhetorical and terminological conventions of shipping forecasts. Several outputs follow the canonical structure of the genre. For instance, Sentence 1 (*Tempête ‘ALEX’ 995 hPa au large de Brest se déplace vers l’est* – “Storm ‘ALEX’ at 995 hPa off Brest is moving east”) mirrors authentic patterns in which storms are tracked by pressure and movement. Sentence 2 (*Front froid 1012 hPa au nord de la Biscaye s’intensifie [...]* – “Cold front at 1012 hPa north of the Bay of Biscay is intensifying”) aligns with standard templates for describing frontal systems (Condamines 2018). Sentence 4, which depicts an anticyclone extending northeastward, also reflects phrasing observed in the corpus, using process and location.

However, several sentences confirm Grok’s lack of meteorological ontology. In Sentence 6, it incorrectly attributes a movement predicate to *houle* (“swell”), which, in expert forecasts, denotes a state rather than an active agent, another example of frame mismatch in the sense of frame semantics (Fillmore 1976). Similarly, Sentence 3 deviates from domain-specific conceptual patterns by depicting a Genoa low deepening westward, a trajectory inconsistent with observed meteorological behavior. In Sentence 8, Grok states that a *dépression* (“low-pressure system”) *fusionne* (“merges”) with a front, an event unattested

in the corpus and semantically incompatible with Frame-Based Terminology, which holds that specialized units encode non-interchangeable conceptual frames (Faber 2015; L’Homme 2020). These deviations confirm that Grok reproduces phraseological structures without integrating expert knowledge.

Grok reproduces the formal and phraseological patterns of marine forecasts, but it lacks the epistemic precision needed to align lexical choices with the underlying conceptual frames of meteorology. As a result, it generates sequences that are syntactically plausible yet conceptually incorrect, revealing the gap between surface genre conformity and true terminological adequacy. In controlled-language contexts, such errors may have significant safety implications.

5.2.3 Comparison between Claude and Grok

To summarize the observations, this section compares the types of errors produced by Claude and Grok. While both models successfully reproduce the surface phraseology characteristic of the genre, their outputs diverge substantially in terms of conceptual accuracy, frame compatibility, and adherence to domain-specific constraints. The following comparison highlights the limitations of each system.

Error Type	Claude-opus-4	Grok-4-Fast
Frame Mismatch	Frequent: assigns pressure values to <i>thalweg</i> , uses dynamic predicates for non-dynamic entities (e.g. <i>dorsale s’oriente</i> ; ridge shifts)	Frequent: treats <i>houle (swell)</i> as an agent, describes impossible trajectories (e.g. Genoa low deepening westward)
Conceptual Error	Moderate: introduces non-attested pressure variations; confuses semantic roles associated with meteorological categories	High: invents implausible events such as a <i>dépression</i> merging with a front; overpopulates semantic roles
Register Error	Occasional: uses generic geographical terms instead of nautical zones (e.g., <i>Méditerranée occidentale</i>)	Rare: mostly stays within genre-appropriate lexicon but may drift into verbose or overly narrative phrasing
General Trend	Strong control of terminology and structure but struggles with ontological consistency	Strong rhetorical fluency but conceptually less reliable

Figure 7: Error types produced by Claude-opus-4 and Grok-4-fast

Collectively, these errors point to different weaknesses in Claude and Grok: Each model exhibits a distinct balance between structural fidelity and conceptual robustness. This supports the idea that apparent genre mastery does not necessarily entail epistemic adequacy. These findings pave the way for the concluding discussion, where we reflect on the broader implications of these limitations and outline future directions for improving LLM performance in highly constrained, ontology-dependent domains.

6 Conclusion, limits and outlook

This study examined the capacity of LLMs to reproduce the lexico-grammatical regularities and discursive organization of French shipping forecasts under zero-shot conditions. Although both models demonstrated an impressive ability to replicate the surface mold of forecast discourse, including its structural segmentation, standardized phraseology, and syntactic compression, they also exhibited systematic conceptual limitations. Claude adhered more closely to terminological constraints but struggled with frame compatibility, whereas Grok produced rhetorically more fluent outputs yet displayed a tendency to misassign semantic roles and overpopulate predicates. In both cases, the lack of an internalized meteorological ontology resulted in conceptually inconsistent or implausible outputs despite surface fluency.

Several limitations temper the generalizability of these findings. First, the number of outputs generated was limited, and the analysis focused on only two model architectures; a broader sample and few-shot settings might increase the performance of the models. Second, the evaluation was conducted without fine-tuning or explicit semantic constraints, meaning that model performance reflects baseline general-purpose behavior rather than domain-specific optimization. Third, although grounded in empirical corpus data, the evaluation remains interpretative and dependent on linguistic, terminological, and conceptual expertise.

Building on these results, future research may examine semantic frame-based prompting as a means of guiding model generation with structured

knowledge representations. Integrating FrameNet-style annotation into prompts could help align lexical items with appropriate conceptual roles, reducing frame mismatches and enhancing epistemic validity. A final direction involves training smaller, domain-adapted models fine-tuned on annotated meteorological corpora, which may outperform general-purpose LLMs in specialized generation tasks while offering greater interpretability and operational control.

References

- AUSSENAC-GILLES, N. / CONDAMINES, A. (2012): "Variation and Semantic Relation Interpretation: Linguistic and Processing Issues." *Conference on Terminology and Knowledge Engineering (TKE 2012)*. Eds. Guadalupe Aguado de Cea / Mari Carmen Suárez-Figueroa / Raúl García-Castro / Elena Montiel-Ponsoda. Madrid: Universidad Politécnica de Madrid. 106–122.
- BENDER, E. M. / GEBRU, T. / McMILLAN-MAJOR, A. / SHMITCHELL, S. (2021): "On the Dangers of Stochastic Parrots: Can Language Models be too Big?" *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT'21)*. Ed. Association for Computing Machinery. New York: ACM. 610–623.
- BOCANEGRA-VALLE, A. (2011): "The Language of Seafaring: Standardized Conventions and Discursive Features in Speech Communications." *International Journal of English Studies* 11.1: 35–56. DOI: 10.6018/ijes/2011/1/137091.
- BOCANEGRA-VALLE, A. (2015): "The Constructing of Shipping Forecasts in English: A Pilot Study." *Procedia – Social and Behavioral Sciences* 212: 278–282. DOI: 10.1016/j.sbspro.2015.11.349.
- CONDAMINES, A. (2018): "For the Development of Ergonomic Linguistics: The Example of Controlled Languages." *Le Travail Humain* 81.3: 205–226. DOI: 10.3917/th.813.0205.
- DONG, Q. / LI, L. / DAI, D. / ZHENG, C. / MA, J. / LI, R. / XIA, H. / XU, J. / WU, Z. / CHANG, B. / SUN, X. / LI, L. / SUI, Z. (2024): "A Survey on In-Context Learning." *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing (EMNLP 2024)*. Ed. Association for Computational Linguistics. Miami, FL: ACL. 1107–1128.

- FABER, P. (2015): "Frames as a Framework for Terminology." *Handbook of Terminology*, Vol. 1. Eds. H. Kockaert / F. Steurs. Amsterdam/Philadelphia: Benjamins. 14–33.
- FILLMORE, C. J. (1976): "Frame Semantics and the Nature of Language." *Annals of the New York Academy of Sciences* 280.1: 20–32. DOI: 10.1111/j.1749-6632.1976.tb25467.x.
- FILLMORE, C. J. (2006): "Frame Semantics." *Cognitive Linguistics: Basic Readings*. Ed. Dirk Geeraerts. Berlin: De Gruyter. 373–400. DOI: 10.1515/9783110199901.373.
- FILLMORE, C. J. / BAKER, C. (2012): "A Frames Approach to Semantic Analysis." *The Oxford Handbook of Linguistic Analysis*. Eds. B. Heine / H. Narrog. Oxford: Oxford University Press. 313–340. DOI: 10.1093/oxfordhb/9780199544004.013.0013.
- GAUTIER, L. (2008): "Fach, Fachsprache und Fachtextsorte: ein 'magisches Dreieck' in der Fachsprachenvermittlung?" *Wort und Text. Lexikologische und textsyntaktische Studien zum Deutschen und Französischen*. Eds. D. Baudot / M. Kauffer. Tübingen: Stauffenburg. 3–13.
- GAUTIER, L. (2014): "Semantisches Prädikat, prädikative Nomina und Argument-linking." *Nomina und Verba im Zusammenspiel*. Eds. J.-F. Marillier / H. Quintin. Tübingen: Stauffenburg. 69–83.
- JOHNSON, B. (1995): "Some Features of Maritime Telex Service Communication." *English for Specific Purposes* 14.2: 127–136. DOI: 10.1016/0889-4906(95)00008-F.
- KILGARRIFF, A. / RYCHLÝ, P. / SMRŽ, P. / TUGWELL, D. (2004): "The Sketch Engine." *Proceedings of the 11th EURALEX International Congress*. Eds. G. Williams / S. Vesier. Lorient: Université de Bretagne Sud. 105–116.
- L'HOMME, M.-C. (2020): *Lexical Semantics for Terminology: An Introduction*. Amsterdam: Benjamins. DOI: 10.1075/tlrp.20.
- LIÉGEOIS, V. / MODICOM, P.-Y. / BONNE, L. (2023): "Construction Grammar in Domain-Specific Discourse." *Constructional Approaches to Nordic Languages*. Eds. E. Coussé et al. Amsterdam: Benjamins. 145–178. DOI: 10.1075/cal.37.06lie.
- LIÉGEOIS, V. / GAUTIER, L. (2025a): "Existential Structures as Specialised Constructions." *ELAD-SILDA* 11 s.p. <https://publications-prairial.fr/elad-silda/index.php?id=1673> (4.12.2025).
- LIÉGEOIS, V. / GAUTIER, L. (2025b): "Domain-Sensitive Weather Phraseology from a Constructionist Perspective: A Case Study Based on Dutch Dialects in Flanders." *Dutch Journal of Applied Linguistics* 14: 1–23. DOI: 10.51751/dujal12632.

- ORNA-MONTESINOS, C. (2012): *Constructing Professional Discourse: A Multiperspective Approach to Domain-Specific Discourses*. Newcastle-upon-Tyne, UK: Cambridge Scholars Publishing.
- PRITCHARD, B. / KALOGJERA, D. (2000): "On Some Features of Conversation in Maritime VHF Communication." *Dialogue Analysis VII: Working with Dialogue*. Eds. M. Coulthard / J. Cotterill / F. Rock. Berlin: De Gruyter. 185–196.
- SINCLAIR, J. (1991): *Corpus, Concordance, Collocation*. Oxford: Oxford University Press.
- THEROINE, S. / GAUTIER, L. / ASCIONE, G. / CRUZ, C. (2025): "La double opposition terminologie/phraséologie et lexique/grammaire revisitée au prisme d'un discours (très) contraint: les bulletins de météo marine." *Synergie Tunisie*. (forthcoming)

KI-Berichterstattung zwischen Euphorie und Realität

KI-Newsletter als Informationsmedium

Abstract

Since its inception in early 2022, Artificial Intelligence (AI) has been debated controversially in the media. The emergence of chatbots, Large Language Models (LLMs), generative AI and AI agents has spurred the discussion about AI strengths, weaknesses and the risks of using AI in business, education, and media. There is rising social concern about the effects of man-machine interaction as well as deep fakes. Therefore, AI media coverage spans from broad acceptance to the critical evaluation of ethical risks. This article examines the perception of AI based upon a qualitative corpus analysis of digital newsletters. The study considers German and English newsletter communications on AI and considers the role of this genre in adapting this disruptive new technology by describing macrostructure, visual elements, style and interactivity with the readers.

Keywords KI in den Medien, KI-Newsletter, KI-Wahrnehmung in Deutschland

1 Einleitung

Seit der Einführung von Chatbots und Large Language Models (LLMs) ist Künstliche Intelligenz (KI) zu einem „Dauerbrenner“ in den digitalen Medien gereift. Das anfängliche Erstaunen und eine gewisse Euphorie über die ver-

meintlich grenzenlosen Möglichkeiten von KI sind einer wachsenden Skepsis und Diskussion über Nutzen und Risiken gewichen. Insgesamt kann man in Deutschland von einer deutlichen Polarisierung in den Einstellungen zu KI ausgehen (Wintergerst 2024; Gansser/Reich 2025; Lammar et al. 2025; Reiss et al. 2025; Thajudeen 2025). Der mediale KI-Diskurs reicht von sachlicher Information zu emotionalen Disputen über KI-Akzeptanz und KI-Ethik. Einerseits bieten KI-Technologien vielfältige Einsatzmöglichkeiten im privaten und geschäftlichen Umfeld. Andererseits wächst die Skepsis gegenüber KI, vor allem vor dem Hintergrund von Informationsmanipulation und nicht klar geregelter Nutzungsbedingungen. Gleichzeitig entstehen Fragen, die die mediale Verantwortung der Berichterstattung über KI und von KI-generierten Inhalten betreffen (Vogler et al. 2023; Afroogh et al. 2024). Allerdings scheint sich die öffentliche Wahrnehmung von KI aktuell von der skizzierten Sichtweise zu einem differenzierten Verständnis zu wandeln, was sich aus der KI-Zugänglichkeit für breitere Nutzerkreise ergibt, aber auch aus der massenwirksam aufklärenden KI-Berichterstattung (Thajudeen 2025).

Der vorliegende Beitrag untersucht anhand von KI-Newslettern, wie die regelmäßige Berichterstattung über die Dynamik der KI-Entwicklung zu einer veränderten Wahrnehmung dieser disruptiven Technologie führt bzw. führen kann. Im Mittelpunkt stehen die Kriterien visuelle Formate von KI, sprachlich-stilistische Mittel und Autoren-Lesenden-Beziehung, die zu einer positiven Meinungsbildung zu KI beitragen können. Zu diesem Zweck wurden von April bis November 2025 je fünf relevante deutsch- und englischsprachige Newsletter zu KI beobachtet und insgesamt 15 Newsletter pro Sprache einer qualitativen Analyse unterzogen, die im Artikel exemplarisch vorgestellt werden.¹ Im Mittelpunkt standen zwei Forschungsfragen:

.....

1 Erste Ergebnisse der Analyse wurden zur 12. Tagung zur Medienlinguistik im September 2025 an der WHZ vorgestellt. Materialien zum Vortrag und Projekt sind abrufbar über <https://padlet.com/busch0lauer/ki-in-den-medien-zwischen-angst-euphorie-und-realit-t-hagt7pyssfa3ic07> (12.12.2025).

- (1) Wie vermitteln KI-Newsletter ein positives Meinungsbild zu KI?
- (2) Welche (kommunikativen) Mittel werden eingesetzt, um die Wahrnehmung von KI beim Adressatenkreis zu beeinflussen?

Der Artikel skizziert zunächst das Meinungsbild über KI in Deutschland, bevor Makrostruktur, Themen und Aufbau der KI-Newsletter sowie exemplarisch Analyseergebnisse vorgestellt werden (Abschnitt 2). Abschnitt 3 diskutiert einige didaktische Möglichkeiten, die Textsorte Newsletter in den Ausbildungsprozess von Studierenden an der WHZ zu integrieren und Abschnitt 4 zieht ein Fazit aus der Untersuchung.

2 Forschungsstand und Parameter der Untersuchung

2.1 Forschungsstand zur Wahrnehmung von KI

KI-Systeme werden von immer mehr Menschen genutzt (MeMo: KI-Meinungsmonitor²). Gut die Hälfte der Bevölkerung hat laut neuen Studien bereits selbst Erfahrungen mit generativer KI gemacht, jüngere Generationen deutlich häufiger als ältere. 67 % der Befragten gaben 2024 an, zumindest gelegentlich KI zu verwenden.³ In der Altersgruppe der 16- bis 19-Jährigen betrug dieser Anteil bereits knapp 96 %. Etwa 63 % der Deutschen beurteilen das Potenzial von KI für gesellschaftlichen Fortschritt positiv, 74 % gaben an, KI als Chance zu sehen, vor allem in Bereichen wie Medizin oder Forschung.

Die YouGov-Befragung mit 1 001 Personen (16.12.2024 bis 2.1.2025) in Social-Media-Posts konnte die Generationenabhängigkeit der KI-Nutzung und der KI-Akzeptanz verifizieren. So lag die KI-Akzeptanz in Posts der Genera-

.....

2 Der Meinungsmonitor KI hat die Nutzung und Einstellung zu KI anhand der Bevölkerungsmeinung und Berichterstattung über KI im Zeitraum 2020–2024 untersucht und ausgewertet.

3 <https://leibniz-hbi.de/hbi-publications/nutzung-und-wahrnehmung-generativer-ki-zur-informationssuche-in-deutschland/> (22.11.2025).

tion der *Millennials* mit 20 % am höchsten, wohingegen die geringste Akzeptanz bei der Gruppe der *Babyboomer* (mit 7 %) bestand.⁴ Allerdings herrscht auch erhebliche Skepsis: 48 % vertrauen nicht in die Korrektheit der KI-Ergebnisse, 43 % geben an, dass sie nicht mit dem KI-generierten Ergebnis zufrieden sind, und 31 % konstatieren sogar falsche Ergebnisse. Weitere Gründe für eine ablehnende Haltung zu KI sind Datenschutzbedenken, die Furcht vor Jobverlust und vor Verlust menschlicher Analysekompetenz. Deshalb ist es ein zentrales Anliegen, die KI-Anwendung zu regulieren und mehr Transparenz zu ihrer Nutzung zu schaffen. Dies gilt insbesondere auch für den Bildungsbereich, in dem Lernende (SchülerInnen und Studierende) KI bereits massiv und oft unreflektiert nutzen und Lehrende sich nur sehr zögerlich mit der neuen Technologie auseinandersetzen, was insgesamt zu erheblichen Dissonanzen im Bildungssektor führen kann.

Laut einer Marktforschungsstudie⁵ wird KI im privaten Bereich am häufigsten für die Korrektur von Texten (50 %), für allgemeine Recherchen (48 %) und für Unterhaltung (44 %) genutzt. 36 % verwenden KI für kreative Textprozesse, 34 % zur Problemlösung, 24 % zur Erstellung und Bearbeitung von Bildern bzw. Videos und nur 20 % für Übersetzungen.

Aus diesen wenigen Studien lässt sich konstatieren, dass das Verhältnis gegenüber der neuen Technologie zwiespältig ist. Obwohl man die Chancen durchaus anerkennt, bleiben Vorbehalte. Aus diesem Grund scheint es notwendig, die meinungsbildende Rolle von Medienformaten, u. a. von Newslettern zu untersuchen, um mehr Vertrauen durch Transparenz von Informationen aufzubauen (vgl. auch van Oorschot/Fucker 2022).

.....

4 <https://www.deutschlands-marktforscher.de/studie-zur-nutzung-von-kuenstlicher-intelligenz-in-deutschland/> (22.11.2025).

5 <https://www.deutschlands-marktforscher.de/studie-zur-nutzung-von-kuenstlicher-intelligenz-in-deutschland/> (22.11.2025).

2.2 Materialkorpus und exemplarische Ergebnisse der Untersuchung

Ein Newsletter ist ein „regelmäßig erscheinendes digitales Informationsmedium, das Abonnenten oder Mitgliedern einer bestimmten Gruppe oder Organisation zugesandt wird. Er dient dazu, die Empfänger über aktuelle Entwicklungen, Neuigkeiten, Angebote oder andere relevante Themen zu informieren.“⁶ Diese digitale Textsorte ermöglicht eine aktuelle, regelmäßige und adressatengerechte Verbreitung von Neuigkeiten für einen fachlich engen (Expertenkreis), aber auch für einen breiten Leserkreis (Laienkreis).

2.2.1 Arten, Inhalte und Aufbau von Newslettern zu KI

Newsletter über KI (und KI-generierte Newsletter zum Thema) unterliegen der starken thematischen Dynamik und Themenvielfalt des Gegenstandes. Seit 2022 haben sich zahlreiche Konkurrenzangebote von Newslettern entwickelt, die erheblich um die Gunst der Lesenden ringen.⁷ Die Textsorte ist geprägt von Multimodalität (indem Text, Bild- und Videoelemente zusammenfließen) und weist einen hohen Grad an Intertextualität auf, was sich in Hyperlinks und Querverweisen zeigt. Die Makrostruktur der Newsletter folgt einer starken Normierung, was die Auswahl und Abfolge von Rubriken betrifft. Damit Inhalte nicht zu textlastig sind, werden zur Auflockerung wiedererkennbare Bildelemente und Logos eingesetzt. Der Wiedererkennungswert von Rubriken erzeugt Verlässlichkeit und Vertrauen, wodurch auch die Einstellung gegenüber Inhalten beeinflusst werden kann.

Die Newsletter lassen sich nach Zielgruppe (allgemein Interessierte, Fachpublikum), thematischem Fokus, Format (KI-Überblick oder Deep Dive,

6 <https://www.cleverreach.com/de-de/push-magazin/email-marketing-strategie-tipps/definitionen-e-mail-marketing/newsletter-definition-was-ist-ein-newsletter/> (22.11.2025)

7 Übersichten zu deutschen Newslettern sind u. a. unter <https://publuu.com/de/knowledge-base/ki-newsletter/> (22.11.2025); <https://www.basichinking.de/blog/2025/08/05/beste-ki-newsletter/> (22.11.2025); <https://www.infinitas.de/blog/die-besten-newsletter-ki-kuenstliche-intelligenz> (22.11.2025); Übersichten zu englischsprachigen Newslettern sind u. a. unter <https://www.letterpal.io/blog/best-ai-newsletters> (22.11.2025); <https://github.com/alternbits/awesome-ai-newsletters> (22.11.2025); <https://www.superhuman.ai/c/best-ai-newsletters> (22.11.2025).

Links, Tipps) und Erscheinungsweise (täglich, wöchentlich, Event-bezogen) klassifizieren. Die Korpusauswahl für die Studie orientierte sich an KI-News-
lettern mit hohen Abonnentenzahlen und einem allgemein an KI interessierten
Adressatenkreis, z. B. *KI-Briefing*, *The Algorithm* sowie KI-Newslettern für spe-
zielle Nutzergruppen (*KI-Campus*, *Das Gelbe vom Ei*). Tab. 1 und 2 geben ei-
nen Einblick in das Korpus und den Charakter der betrachteten KI-Newsletter.

Tab. 1: Prominente deutsche Newsletter zu KI⁸

Name	Themen	Zielgruppe	Erscheinungs- weise	Format
<i>KI-Briefing</i>	Aktuelle KI- Trends und Initiativen	Wirtschaft/ Führungs- kräfte	Wöchentlich	Kurze, welt- weite Berichte über KI- Trends
<i>KI-Campus</i>	Newsletter, Kurse, Videos, Tools	Wissenschaft, Forschung	Regelmäßig	Kurztex- te, Kurse, Videos, Podcasts, Tools
<i>AINauten</i>	KI-News & Tools, Hacks	breite Öffentlichkeit und Geeks	Mehrmals/ Woche verfügbar über Website	Kuratierte, kompakte News, Tipps & Hacks
<i>Applied AI initiative</i>	KI-Beratung für Unterneh- men	Unternehmen	2 x monatlich	News, Fall- beispiele für KI-Nutzung in Unternehmen Trainingskurse
<i>Das Gelbe vom AI</i> (auf LinkedIn)	Branchen- fokus	Unternehmen, Entscheider	Regelmäßig	Podcast-basier- ter Newsletter

Bezogen auf die Inhalte, informieren KI-Newsletter regelmäßig über aktuelle
Trends, neue KI-Tools und Anwendungsbeispiele (*Use Cases*). Häufig werden

.....
8 Zur Bewertung und Beschreibung der KI-Newsletter vgl. <https://blog.hubspot.de/marketing/ki-newsletter> (22.11.2025).

technische Hintergründe, *Best Practices* und die Auswirkungen auf Wirtschaft und Gesellschaft beleuchtet.

Die deutschsprachigen Newsletter fokussieren oft auf Entwicklungen in den deutschsprachigen Ländern (D-A-CH). Sie stellen in einem eher nüchtern gehaltenen Sprachstil sehr praxisnah neueste KI-Entwicklungen vor und verweisen auf Anwendungen im Marketing. Eine Interaktivität mit den Lesenden entsteht implizit durch Aufforderungen, KI-Produkte direkt zu testen und Feedback zur Nützlichkeit der Newsletter-Inhalte zu geben. Richtet sich der Newsletter an einen technikaffinen Adressatenkreis werden auch Fachjargonismen genutzt, der Sprachstil ist dann durch umgangs- und berufssprachliche Elemente geprägt.

Die englischsprachigen Newsletter verfolgen thematisch eine globale Perspektive und gehen stärker auf wirtschaftliche Aspekte von KI ein. Der Fokus liegt auch stärker auf Entwicklungen in Forschung und Entwicklung in der Tech-Szene (Rolle von *Global-Player*-Unternehmen und *Start-ups*).⁹ Die Berichterstattung ist oft stärker geprägt von positiven Emotionen, was sich durch journalistische Berichterstattung mit Storytelling, fiktiven Dialogen, Erfahrungsberichten sowie stilistisch durch Buzzwords, schmückende Epitheta, sprachliche Bilder und Vergleiche zeigt (vgl. Abschnitt 3).

Tab. 2: Prominente englischsprachige Newsletter zu KI

Name	Themen	Zielgruppe	Erscheinen	Format
<i>The Algorithm</i> (MIT)	Neuestes Trends in KI	MIT-Interessierte	Wöchentlich	Dialogform, Storytelling, Verlinkungen
<i>AI Weekly</i>	Breite Übersicht zu KI-News & Tools, Forschung, globale Ereignisse	Fachpublikum, Entwickler	Wöchentlich	Internationaler Überblick, Forschung, Tools

.....
 9 Zur Charakterisierung der KI-Newsletter vgl. <https://ki-trainingszentrum.com/ki-newsletter-2/> (22.11.2025).

Name	Themen	Zielgruppe	Erscheinen	Format
<i>TLDR-AI</i> (steht für „Too Long, Didn’t Read“)	Technische Fokussierung, Code, Tools, Produkte um KI	Tech-Enthusiasten Data Scientists	Täglich	Kurze Analysen zu KI-Innovationen, Machine Learning, Data Science
<i>Generative AI Daily Newsletter</i>	Neueste KI-Nachrichten	Start-ups, Techies	Täglich	Kuratierte Inhalte ohne Werbung
<i>Super-human</i> von Tech-Spezialist Zain Khan	Tools & Trends	Experten, Innovatoren	Täglich	Tool-Vorstellungen, Social-Media-Trends, lustige Prompts

Beispielhaft seien an dieser Stelle die Makrostruktur und Rubriken je eines deutsch- und eines englischsprachigen Newsletters vorgestellt. Beiden gemeinsam ist die Verschmelzung von Briefelementen und journalistischer Berichterstattung. Allerdings unterscheiden sie sich im Stil und der Ansprache der Lesenden: im Deutschen mehrheitlich neutral-sachlich, im Englischen eher emotionaler.

KI-Briefing (<https://ki-briefing.com/>) ist der größte deutschsprachige KI-Newsletter für Führungskräfte (60 000 Abonnenten und 120 000 Follower auf LinkedIn)¹⁰. Nach neutraler Begrüßung der Lesenden (*Guten Morgen!*) werden die Ereignisse der Woche kurz in 3 textuell ausformulierten Absätzen zusammengefasst, wobei der letzte Absatz stets zur Interaktion aufruft und das Feedback der Lesenden abfragt (*„Teilen Sie uns gerne mit, welche Impulse für Sie am wertvollsten sind – und helfen Sie dabei, unser Briefing in Ihrem Netzwerk bekannt zu machen. Gemeinsam können wir eine Community aufbauen, die voneinander lernt und sich gegenseitig stärkt.“*). Unter der Rubrik *Was Sie in diesem Briefing erwartet* folgen die verlinkten Kategorien *News*, *Deep Dives*,

10 Beispielnewsletter unter: <https://ki-briefing.com/p/draghi-reformen-versanden-im-bru-sseler-verwaltungsapparat> (22.11.2025).

In aller Kürze, Videos & Artikel, Impuls, Umfrage, Meinung der Redaktion, Praxisbeispiel (Fallbeispiel) und *YouTube* (Entwicklungen auf der Plattform) sowie *Werben im KI-Briefing* und *Vorschau* (*Und nächste Woche ...*). Diese Kategorien ermöglichen ein selektives Lesen nach persönlichen Präferenzen. Jede Kategorie wird durch ein zum Text passendes Bild (Fotos, reale, abstrakte und KI-generierte Bilder) eingeleitet und durch Zwischenüberschriften gut strukturiert. Auch hier zeigen sich klar wiederkehrende Unterkategorien von Texten (*Zusammenfassung, Kernaussagen, Warum das wichtig ist*). Am Ende ist das Feedback der Lesenden durch Emoticons (*Es war super, Es war in Ordnung, Es war schlecht*) gefordert. Alle KI-Briefings sind über die Homepage als Archiv abruf- und so auch wiederauffindbar.

The Algorithm ist der montäglich erscheinende Newsletter des renommierten *MIT Technology Review*-Magazins (aktuell mit der *Financial Times*, FT) und richtet sich an ein interessiertes Fachpublikum und Spezialisten, die tiefergehende Analysen zu KI-Themen wünschen. Der Aufbau dieses Newsletters unterscheidet sich von deutschsprachigen Newslettern durch einen Aufhänger in Form eines Storytellings (meist von Senior Editors von MIT), über einen Dialog (Vertreter MIT und FT-Reporter) oder einen längeren Leitartikel (*Deep Dive*). Thema bzw. Akteure werden in einem einleitenden Abschnitt nach „*Welcome back to the State of AI*“ vorgestellt. Die Texte werden untersetzt durch Bildcollagen, die die Akteure/Dialogpartner zeigen oder deren Bezug zu KI. Dieser Newsletter eignet sich nicht für das schnelle Screening von Informationen, sondern für *Deep Dives*. Die Rubriken sind sehr viel stärker mit abstrakten bzw. KI-generierten Bildern als Aufmerksamkeits- und Leseanreiz illustriert. Außerdem werden viele Hyperlinks verwendet und interaktive Galerien, Videos und Hotspots zu KI-Tools integriert, was die Inhalte anschaulicher macht. Auch in diesem Newsletter werden Fragen nach der Bedeutung von KI im Kontext von Wirtschaft, Politik und Ethik gestellt. Unter *Bits and Bytes* gibt der Newsletter kurze, praktische Tipps zu KI-Anwendungen und Tools sowie deren Evaluation für die Bereiche Marketing, Content-Erstellung, Web-Design und Automatisierung. Jeder Abschnitt liefert weiterführende Links zu MIT-Stories, um das jeweilige Thema noch weiter zu vertiefen. Werbetekhnisch wirksam sind die Querverweise zu weiteren Newslettern des MIT.

Insgesamt ist festzustellen, dass die Makrostruktur der KI-Newsletter stark durch Rubriken normiert ist, die eine informationsverdichtende und strukturierende Funktion haben. Die Verschmelzung von Text und visuellem Material sowie Querverweisen durch Hyperlinks eröffnet die Abstufung der Berichterstattung und eine adressatengerechte Information. Die allen Newslettern eigene Briefform personalisiert die Texte und evoziert beim Lesenden die individuelle Zugehörigkeit zu einer Newsletter-Gemeinschaft. Interaktive Elemente (Aufforderungen zum Testen von Tools und Feedback) lassen die Lesenden Teil der Diskussion sein, bauen Vertrauen und eine positive Wahrnehmung von Inhalten auf.

2.2.2 Visualisierung von KI in den Newslettern

Die bildliche bzw. metaphorische Darstellung von KI und ihre Wirkung auf die Wahrnehmung der Technologie hat sich zu einem neuen Forschungsfeld in der Sprach- und Kommunikationswissenschaft entwickelt (Leidecker-Sandmann et al. 2025; The Royal Society¹¹). Wallenborn (2022) hat umfassend zur Metaphorisierung von KI geforscht. Da (bewegte) Bilder einer Sache eine prägende positive oder negative Wahrnehmung bzw. Einstellung zu dieser Sache evozieren können, konnte sie nachweisen, dass anthropomorphisierte Darstellungen von KI einerseits positiv wirken, weil ein „menschenähnliches Etwas“ zu besserer Arbeitsleistung verhilft, aber andererseits auch negative Wahrnehmungen auslöst, da sich Menschen von KI „verdrängt“ oder „überfordert“ fühlen.

Die Newsletter weisen einen unterschiedlichen Grad an bildlicher Darstellung von KI auf. Während eher akademisch ausgerichtete Newsletter (*KI Campus*, *KI Briefing*, *The Algorithm*) weitgehend auf (KI-generierte) Bilder verzichten und eher Symbolik oder abstrakte Grafiken nutzen, verwenden Newsletter, die stärker an die allgemeine Öffentlichkeit bzw. eine Community von KI-Interessierten gerichtet sind, ein Repertoire an Bilddarstellungen, die oft KI-generiert sind (menschenähnliche Roboter in Männer- oder Frauengestalt, futuristische humanoide Gesichter, Arme und Hände), digitale

.....
11 <https://royalsociety.org/-/media/policy/projects/ai-narratives/AI-narratives-workshop-findings.pdf> (2018), (12.12.2025).

Avatare, Abbildungen, in denen es zur Verschmelzung von Mensch und Maschine kommt. Eher neutral sind Abbildungen mit neuronalen Netzen, die „denkende“, „lernende“ und „vernetzte“ Systeme beschreiben. Abbildungen von Platinen, abstrakten Datenströmen, Zahnrädern oder stilisierten Computer-Systemen betonen die technische Seite von KI als Algorithmus.¹² Insgesamt scheint die visuelle Darstellung von KI-Themen aktuell stärker durch abstrakte Bilder und Schemata geprägt zu werden, da die Visualisierung von KI (Algorithmus) durch Anthropomorphisierung in zwei Wahrnehmungsrichtungen wirken kann: Einerseits wird damit positiv Nähe und Verbundenheit zwischen Menschen und Maschine ausgedrückt, andererseits schürt der Vergleich von KI mit menschlichen Eigenschaften auch Bedenken, dass ein „menschenähnliches Wesen“ die Kontrolle über Prozesse übernimmt. Deshalb werden wieder reale Bilder (statt KI-generierter Bilder) und Icons als Symbole für Digitalisierung (binäre Codes, Schaltkreismuster und futuristische Landschaften und Grafiken mit Datenströmen) verwendet, da sie negative Emotionen eher ausschließen.

2.2.3 Sprachliche Gestaltung von Newslettern

Die sprachliche Gestaltung der KI-Newsletter unterstreicht das Bestreben der Autorinnen und Autoren, verständlich über Innovation und Akzeptanz von KI in verschiedenen Anwendungsfeldern zu berichten. Die Mischform aus Brief und journalistischer Berichterstattung sowie direkter Feedback-Option führt zu Nähe mit den Lesenden als Teil eines *Ecosystems* bzw. einer Community.

Typisch für die Newsletter sind die individuelle namentliche Ansprache des Abonnenten oder einer Gruppe von Adressaten (z. B. *AINauten*, <https://www.ainauten.com/>: *Moin AINauten, Willkommen zur neuen Ausgabe deines Lieblingsnewsletters!*). Im Deutschen ist auffällig, dass institutionelle Newsletter die Anredeform „Sie“ nutzen (*KI-Briefing: Was Sie in diesem Briefing erwartet*), wohingegen die Tech-Szene (z. B. *AINauten*) die Anredeform „Du“ präferiert (*AINauten: Wir filtern für dich den ganzen KI-Hype und liefern dir*

.....

12 Beispiele für Abbildung von KI sind im Vortrag zur Medienlinguistik (vgl. Padlet: <https://padlet.com/busch0lauer/ki-in-den-medien-zwischen-angst-euphorie-und-realit-t-hagt7pyssfa3ic07>).

3x pro Woche die relevantesten News, Tools und Hacks – ohne Tech-Blabla, dafür mit echtem Mehrwert für deinen Arbeitsalltag). Dankesbekundungen für das Lesen der Beiträge und der Bezug auf frühere Interaktionen mit den Lesenden drücken Wertschätzung in der Kommunikation aus. Diese Höflichkeitsprinzipien aus englischsprachigen Newslettern waren in einigen Fällen prägendes Vorbild für deutsche Newsletter.

Neben der visuellen Darstellung und Interaktivität mit den Lesenden sind journalistisch geprägte Kurztexte zur Strukturierung der Rubriken markant. Die Headlines sind meist mit Hyperlinks unterlegt und führen zu weiteren Informationsquellen. Interessant ist, dass sich die Überschriften und Teilüberschriften im Deutschen und Englischen durchaus in Wortwahl, Stil und Ton unterscheiden. Während deutsche (akademisch ausgerichtete) Newsletter eher nüchtern und informativ bleiben (z. B. bei *KI-Campus: KI-Kompetenzen für die Hochschulverwaltung: Verantwortungsvoller Umgang in der Praxis; Zukunftsimpulse: Allianz für Future Skills stärkt KI-Bildung an Hochschulen; KI im Überblick – Was Sie wissen müssen*), sind die Headlines in englischsprachigen Newslettern emotionaler und durch Stilmittel wie konnotierte Lexik, Metaphorik, Alliteration, Reime etc. geprägt (z. B. *The Algorithm: Can an AI doppelgänger help me do my job?; Superhuman: AI adoption eats entry-level jobs; aiweekly.co: Bubble Trouble: AI rally shows cracks as investors question risks; AI startup mania: 25-year-olds become Billionaires while Bezos plots his comeback, Supercharge your workflow with AI; AI news that matters*).

Im deutschen Newsletter-Kontext steht Sachlichkeit im Vordergrund, während englischsprachige Newsletter eine aktivierende und trendorientierte Sprache nutzen. Die Inhalte werden fachlich auf die Bedürfnisse der Lesenden zugeschnitten. Die Textgestaltung ist optisch gut strukturiert, angepasstes Fachvokabular und Alltagssprache werden genutzt, um Inhalte ansprechend und verständlich zu gestalten. Eine Überfrachtung mit technischen Inhalten oder Bezeichnungen wird zugunsten von Beispielen oder durch praktische Nutzungsszenarien meist vermieden. Emotionale Formulierungen und die Verwendung von Ausdrücken, die Neugier, Optimismus und Inspiration vermitteln, führen zu einer positiven Gesamtstimmung von Texten. Zugleich wird durch die Einladung zur Interaktion das Gemeinschaftsgefühl gestärkt (Bei-

spiele: *Entdecken Sie die Zukunft; Gemeinsam mehr erreichen, Was denken Sie? Probieren Sie es aus!*).

Ähnlich wie im Marketing wird in den englischsprachigen Texten *Storytelling* eingesetzt. Kurze Geschichten, Anekdoten oder reale Anwendungsfälle machen KI erfahrbar und setzen auch emotionale Anknüpfungspunkte, die wichtig sind für die Ausbildung von Überzeugungskraft. Besonders wirksam ist der Rückgriff auf die Erfahrungswelt der Lesenden (vgl. Beispiel 1 und 2). Berichte, die KI als innovative und hilfreiche Technologie darstellen, führen dazu, dass sie mit Modernität und Fortschritt assoziiert werden (vgl. Beispiel 3). Beispiel 4 verdeutlicht aber auch den persönlichen Gestaltungsraum von Newslettern. AInauten fokussiert besonders auf die Anwendung von Tools und nutzt eine sehr familiäre Umgangssprache.

Beispiel 1 aus *The Algorithm* (7.7.2025)

Hi, Algorithm readers! I'm Grace Huckins, MIT Technology Review's newest AI reporter, and I'm excited to be writing The Algorithm this week. I'm definitely an all-around AI geek—I don't think I would have landed this gig otherwise!—but I'm particularly obsessed with the use of AI in scientific research. (...)

Today's AI landscape is defined by the ways in which neural networks are unlike human brains. A toddler learns how to communicate effectively with only a thousand calories a day and regular conversation; meanwhile, tech companies are reopening nuclear power plants, polluting marginalized communities, and pirating terabytes of books in order to train and run their LLMs.

*But neural networks are, after all, neural—they're inspired by brains. Despite their vastly different appetites for energy and data, large language models and human brains do share a good deal in common. They're both made up of millions of subcomponents: biological neurons in the case of the brain, simulated "neurons" in the case of networks. They're the only two things on Earth that can fluently and flexibly produce language. **And scientists barely understand how either of them works.***

Beispiel 2 aus *The Algorithm* (25.8.2025)

Welcome back to The Algorithm!

Stop me if you've heard this one before.

The AI learns it is about to be switched off and goes rogue, disobeying commands and threatening its human operators.

It's a well-worn trope in science fiction. We see it in Stanley Kubrick's 1968 movie 2001: A Space Odyssey. It's the premise of the Terminator series, in which Skynet triggers a nuclear holocaust to stop scientists from shutting it down.

Those sci-fi roots go deep. AI doomerism, the idea that this technology—specifically its hypothetical upgrades, artificial general intelligence and super-intelligence—will crash civilizations, even kill us all, is now riding another wave.

The weird thing is that such fears are now driving much-needed action to regulate AI, even if the justification for that action is a bit bonkers.

Beispiel 3 aus *applied AI initiative* (24.7.2025)

Guten Tag Name der Person,

*Das **Agentic AI Meetup** vergangene Woche, gemeinsam mit XXX veranstaltet, war ein energiegeladener Austausch von Ideen, Anwendungsfällen und visionären Ansätzen – alles drehte sich um die nächste Entwicklungsstufe von Enterprise AI: **Agentic Systems**. Diese Agenten analysieren Daten nicht nur – sie handeln, lernen und verändern, wie wir arbeiten. Hier sind die wichtigsten Highlights:*

Beispiel 4 aus *AIonauten* (15.9.2025)

Moin AIonauten,

Willkommen zur neuen Ausgabe deines Lieblingsnewsletters!

Großes Update bei ChatGPT: neu kannst du nicht nur Output generieren lassen, sondern ChatGPT kann auch direkt Dinge in deinen Apps für dich tun! Extrem mächtig, aber es gibt ein paar Dinge dabei zu beachten. Zudem haben wir Replit's Agent 3 getestet und zeigen dir, wie AI dir die besten Life Hacks verrät. Ready?

Das haben wir heute für dich im Gepäck: (...) Los geht's!

Insgesamt sind die deutschsprachigen Newsletter fachlich komplexer, setzen die Kenntnis von Termini voraus, z. B. *Algorithmus*, *Neuronale Netze*, *Large Language Model*. KI wird mit wertenden Epitheta verknüpft (*revolutionär*, *innovativ*, *maßgeschneiderte Anwendung*). Teilweise wird die KI zum handelnden Subjekt, z. B. *Die KI lernt*, *Das System erkennt*. Die englischsprachigen Newsletter nutzen oft Buzzwords, auffällig sind Alliterationen und konnotierte Begriffe (*smart solutions*, *Bringing business benefits*) sowie Slogans (*Boost your productivity with AI*). In der Interaktion mit dem Lesenden werden Aufforderungsverben eingesetzt: *Unlock new efficiency*; *Discover how AI transforms*; *Get insights powered by AI*.

Eine gewisse Gleichförmigkeit in der Makrostruktur sorgt für einen einheitlichen Ton in den Newslettern, der Wiedererkennungswert hat. Storytelling und persönliche (fiktive) Erfahrungsberichte vermitteln Authentizität und einen glaubwürdigen Stil, der in der Folge Vertrauen schafft und Meinungen zu festigen hilft. Der überwiegend neutrale, sachliche Sprachstil der deutschsprachigen Newsletter vermittelt den Lesenden Authentizität und gleichzeitig solide Informationsvermittlung, was in Zeiten von *Fake News* wichtig ist, um Vertrauen und Transparenz zu wahren.

In Zukunft werden Newsletter zu KI zunehmend durch KI-Anwendungen automatisiert generiert werden. Dann können KI-Systeme die spezifischen Präferenzen von Adressaten durch Analyse des Leseverhaltens erkennen und darauf aufbauend personalisierte Leseempfehlungen geben, auch gepaart mit der Option und Gefahr der Manipulation. KI-Agenten werden individualisierte Newsletter automatisch kompilieren. Schon heute sind in dieser Richtung, z. B. über <https://publuu.com>, Möglichkeiten vorhanden. Damit zeigt sich, dass KI sowohl Potenzial als auch Risiko in sich trägt.

3 Didaktisches Potenzial der Textsorte Newsletter (zu KI)

Im Folgenden werden einige Szenarien zur didaktischen Integration der Textsorte Newsletter (zu KI) in die Ausbildung von wirtschafts- und medien-

sprachlich ausgerichteten Studiengängen vorgestellt. An der WHZ betrifft dies die Studiengänge *Wirtschaftskommunikation Deutsch* und *Angewandte Deutsch-Französische Medienkommunikation*, in denen auch projektorientiertes Lernen praktiziert wird. Die teilnehmenden internationalen Studierenden studieren Deutsch als Fremdsprache und bereiten sich im Studium auf eine mehrsprachige Berufspraxis in den Feldern Marketing, Public Relations und Media Research vor, wo KI-Kompetenz (*AI literacy*) in Zukunft rezeptiv und produktiv gefordert wird.

Newsletter zu KI können als Einführung in die Arten und Funktionen von digitalen Textsorten in die Lehrveranstaltungen einbezogen werden und verfolgen zwei Ziele: Einerseits wird das Textsortenwissen der Studierenden zu einer prominenten Textsorte erweitert und andererseits erfahren die Studierenden, die in der Regel wenig über die Dynamik der KI-Entwicklung wissen, mehr über aktuelle Trends, Tools und Anwendungsszenarien von KI. Die Vermittlung der Textsorte sollte zwischen rezeptions- und produktionsorientierten Phasen wechseln, um sowohl textlinguistische als auch medienkritische Kompetenzen aufzubauen. Folgende Abfolge von Szenarien wäre denkbar:

In einem 1. Szenario werden Newsletter zu KI (Deutsch und Englisch) recherchiert und die Studierenden erarbeiten textsortenspezifische Merkmale und medienlinguistische Kategorien (Adressatenbezug, Tonalität, Multimodalität, Hypertextstruktur). Sie entwickeln ein Analyse-Raster aus linguistischer Perspektive zu Makrostruktur, sprachlichen Mustern, Verständlichkeitsstrategien, Framing von KI, Positionierung von Institution und Autorenschaft. In Gruppenarbeit wird je ein Newsletter analysiert und die Ergebnisse werden im Plenum oder in Form einer Poster-Session mit Fokus auf medien- und sprachkritische Aspekte vorgestellt (z. B. *Wie wird KI als Chance/Risiko gerahmt?*). Im Mittelpunkt dieses Szenarios stehen die kritische Reflexion der Rolle von Newslettern in der KI-Öffentlichkeit und die Stärkung der digitalen Recherchekompetenz sowie das Nutzen von typischen Struktur-, Sprach- und Gestaltungsmitteln der Textsorte.

In einem 2. Szenario steht die Leserperspektive im Zentrum. Nachdem sich die Studierenden für das Verfolgen/Abonnieren eines Newsletters entschieden haben, liegt das Augenmerk auf der Lesenden-Perspektive, auf Glaubwürdigkeit des Materials, auf Meinungs- und Vertrauensaufbau. Die Studierenden

erhalten zwei Varianten eines KI-Newsletters (z.B. förmlich vs. dialogisch) und protokollieren, was die Rezeption erleichtert, was Verständnisprobleme evoziert, wie Kernaussagen formuliert sind. Sie leiten daraus Schritte für die Planung und Produktion eines eigenen Newsletters für eine klar abgegrenzte Zielgruppe (z.B. Studierende) ab. Gleichzeitig werden die *Dos and Don'ts* der Textsorte erfasst.

In Szenario 3 beginnt die kleinteilige Produktion von Texten/Textbausteinen eines Newsletters. Da die Studierenden Texte in der Fremdsprache Deutsch verfassen, wird in Gruppen eine Textbaustein-Werkstatt durchgeführt, z.B. bearbeitet eine Gruppe Überschriften/Betreffzeilen, eine weitere Teaser (Verdichtung eines längeren KI-Artikels mit klarer Leserführung) und eine Gruppe formuliert Call-to-Action-Elemente (z.B. *Jetzt anmelden, Mehr erfahren*).

Mit Szenario 4 startet das semesterbegleitende Projekt, bei dem Studierende in Redaktionsgruppen eine Newsletterausgabe konzipieren, schreiben, redigieren und Testversand durchführen. Als autonome Lernende entwickeln sie ein Redaktionskonzept für den Newsletter (Zielgruppe, Kommunikationsziele, Rubriken) und verteilen die Arbeitsressorts (Rollenverteilung, Chefredaktion, Ressorts, Lektorat etc.). In Redaktionssitzungen werden die verfassten Texte vorgestellt und evaluiert sowie durch Interviews oder Statements von Lehrenden/Studierenden weiterentwickelt. Nach Fertigstellung der Texte kann der Newsletter im geschützten Rahmen (z. B. an Studierende der Fakultät) verschickt werden und mit einer Online-Umfrage zu Verständlichkeit und Überzeugungskraft evaluiert werden. Das Feedback wird von den Studierenden erneut bewertet.

In Szenario 5 experimentieren die Studierenden mit KI-gestützter Textarbeit und reflektieren die über Prompts erstellten KI-Texte für einen Newsletter. Als Vergleichsaufgabe lassen die Studierenden einen kurzen KI-Newsletter-Entwurf erzeugen und evaluieren diesen nach journalistischen Qualitätskriterien (Stil, Genauigkeit, Transparenz, Bias). Als Reflexion (z.B. im Portfolio) dokumentieren die Studierenden wie sich die KI-Unterstützung auf Textsorte, Sprache, Ton, Verständlichkeit etc. ausgewirkt hat. Schließlich formulieren sie einige Leitlinien zur Nutzung von KI-Tools bei der Newsletter-Produktion (z. B. Kennzeichnung, Quellenprüfung, Datenschutz).

Im Seminarverlauf sammeln die Studierenden ihre Analysen und Texte in einem Portfolio, das schließlich als Leistungsnachweis dient und benotet wird. Obligatorische Bestandteile dieses Portfolios sind: (1) Analysebeispiel eines Newsletters mit Analyse-Raster (rezeptive Kompetenz); (2) eigener, mehrfach überarbeiteter KI-Newsletter (Vorversion und Peer-Feedback); (3) kurzer Reflexionstext zur Frage „Wie prägt Newsletter-Kommunikation die Wahrnehmung von KI?“.

Jede Seminarsitzung sollte die Phasen Einstieg – Analyse – Produktion – Reflexion beinhalten und rezeptives und produktives Arbeiten mit authentischem Material verbinden. Templates können die Kriterien für die Struktur von Newslettern, die Textproduktion und auch die Bewertung für Beiträge unterstützen. Auf diese Weise kann praxisnah Fachwissen aufgebaut werden und beruflichen Ansprüchen des späteren Berufsbildes besser entsprochen werden. Dieses Vorgehen setzt jedoch bei den Studierenden und Lehrenden Offenheit gegenüber KI voraus, die leider noch nicht in jedem Fall gegeben ist.

4 Die meinungsbildende Funktion von Newslettern zu KI

Die vorläufige qualitative Analyse von Newslettern zu KI zeigt, dass diese digitale Textsorte mit hohen Abonnentenzahlen maßgeblich zur Verbreitung von Informationen über Entwicklungen im Bereich KI beiträgt. Folglich hat die Art und Weise, wie diese Kommunikate aufgebaut und gestaltet sind (persönlich adressierend, informationskondensierend, sprachlich und bildlich ansprechend), einen entscheidenden Einfluss auf die Rezeption, die Wahrnehmung und auch auf die Einstellung zu KI.

Positive Berichterstattung über KI und die Interaktivität mit den Lesenden über Newsletter fördert nicht nur die Informationsweitergabe, sondern auch das Vertrauen in die KI-Möglichkeiten und die KI-Nutzung. Entscheidend ist, dass die Inhalte glaubwürdig und nachvollziehbar dargestellt werden, was im untersuchten Korpus auffallend gut gelungen ist. In deutschen Texten über-

zeugt die nüchterne Sachlichkeit, in den englischsprachigen Texten der Rückgriff auf den emotionalen Erfahrungshorizont der Lesenden.

Durch das Zusammenspiel von visuellen Darstellungen, einem von Fachtermini und Umgangssprache gleichermaßen durchmischten verständlichen Sprachstil können Newsletter helfen, KI-Entwicklungen und Konzepte einem breiten Publikum verständlicher zu machen und so auch die Vorteile und Potenziale der neuen Technologie transparent herauszustellen.

Durch die Methoden des Storytellings und Narrativs, die in englischsprachigen Newslettern genutzt werden (insbesondere *The Algorithm*), entsteht Neugier, KI zu nutzen. Vereinfachte Erklärungen und alltägliche Beispiele helfen, eigentlich komplexe Sachverhalte zu vereinfachen. So können Newsletter durch ihre regelmäßige Erscheinungsweise und Aktualität auch als eine Option zur Popularisierung von Fachwissen betrachtet werden.

Die Kombination aus verständlichem Sprachstil und multimedialer Gestaltung trägt maßgeblich zu Vertrauen, Akzeptanz und Engagement der Adressaten bei. Eine glaubwürdige, authentische und transparente Berichterstattung über KI-Entwicklungen kann helfen, die Vorbehalte gegen die Technologie abzubauen und mehr Menschen zu überzeugen, dass die KI-Ära viele Vorteile für ihr privates und berufliches Leben eröffnet, ohne potenzielle Risiken zu unterschätzen.

Literatur

- AFROOGH, SALEH / ABKARI, ALI / MALONE, EMMIE / KARGAR, MOHAMMADALI / ALAMBEIGI, HANANEH (2024): „Trust in AI: Progress, Challenges, and Directions.“ *Humanities and Social Sciences Communications* 11: 1568. DOI: 10.1057/s41599-024-04044-8.
- GANSSE, OLIVER / REICH, CHRISTINA (2025): *FOM VerhaltensRadar: Einstellungen, Chancen und Risiken der Künstlichen Intelligenz in Deutschland. FOM Sommerstudie 2025 – Ergebnisse*. München: FOM Hochschule: Institut für Empirie und Statistik.
- LAMMAR, D. / HORST, M. / MÜLLER, R. (2025): „AI in the German Media: Narratives of AI-in-Particular and AI-in-General in German Media Reporting About Artificial Intelligence.“ *Digital Journalism* 1–19. DOI: 10.1080/21670811.2025.2493759.

- LEIDECKER-SANDMANN, M. / LÜDERS, T. / MOSER, C. / BOGER, V. R. / LEHMKUHL, M. (2025): „More than Humanoid Robots and Cyborgs? How German Print Media Visualize Articles on Artificial Intelligence.“ *JCOM* 24.02: A09. DOI: 10.22323/2.24020209.
- MeMo: KI-Medienmonitor. *Analyse der Medienberichterstattung zum Thema Künstliche Intelligenz. CAIS Research for the Digital Age. Bochum: Ruhr-Universität Bochum.* <https://www.cais-research.de/forschung/memoki/memoki-medienanalyse/> (18.11.2025).
- REISS, MICHAEL V. / KNOR, EVA LUISE / STÖWING, EZRA / MERTEN, LISA / MÖLLER, JUDITH (2025): *Zwischen Neugier und Skepsis: Nutzung und Wahrnehmung generativer KI zur Informationssuche in Deutschland.* (Arbeitspapiere des Hans-Bredow-Instituts | Projektergebnisse 76). Hamburg: Hans-Bredow-Institut. DOI: 10.21241/ssoar.100907.
- THAJUDEEN, DARWISH (2025): „Wie sich die öffentliche Wahrnehmung von KI im letzten Jahr entwickelt hat.“ <https://de.mi4people.org/post/wie-sich-die-%C3%B6ffentliche-wahrnehmung-von-ki-im-letzten-jahr-entwickelt-hat> (22.11.2025).
- THE ROYAL SOCIETY (2018): *Portrayals and Perceptions of AI and Why They Matter.* London: The Royal Society. <https://royalsociety.org/-/media/policy/projects/ai-narratives/AI-narratives-workshop-findings.pdf> (12.12.2025).
- VAN OORSCHOT, FREDRIKE / FUCKER, SELINA (Hrsg.) (2022): *Framing KI: Narrative, Metaphern und Frames in Debatten über Künstliche Intelligenz.* (FEST Forschung, Band 2). Heidelberg: heiBOOKS. DOI: 10.11588/heibooks.1106.
- VOGLER, DANIEL / EISENEGGER, MARK / FÜRST, SILKE / UDRIS, LINARDS / RYFFEL, QUIRIN / RIVIÈRE, MAUDE / SCHÄFER, MIKE S. (2023): „Künstliche Intelligenz in der journalistischen Nachrichtenproduktion: Wahrnehmung und Akzeptanz in der Schweizer Bevölkerung.“ *Jahrbuch Qualität der Medien 2023.* Hrsg. Forschungszentrum Öffentlichkeit und Gesellschaft. Basel. 33–45. DOI: 10.5167/uzh-235608.
- WALLENBORN, JOHANNA T. (2022): *KI als fliegendes blaues Gehirn? Wie Metaphern unsere Vorstellung von KI beeinflussen.* <https://www.hiig.de/publication/ki-als-fliegendes-blaues-gehirn-wie-metaphern-unsere-vorstellung-von-ki-beeinflussen/> (6.12.2025).
- WINTERGERST, RAL (2024): *Künstliche Intelligenz in Deutschland – Status quo und Ausblick.* Bitkom. <https://www.bitkom.org/sites/main/files/2024-10/241016-bitkom-charts-kuenstliche-intelligenz-final.pdf> (22.11.2025).

FACHSPRACHENVERMITTLUNG MIT KI-UNTERSTÜTZUNG

KI als Schlüssel zur Fachsprache

Personalisierte Lernwege und Kollaboration im Fachunterricht

Abstract

The rapid development of Artificial Intelligence (AI) opens up new didactic possibilities for fostering disciplinary language competence in the natural sciences. AI-supported tools enable personalised access to disciplinary texts, promote self-directed and cooperative learning, and facilitate heterogeneous learning groups' access to specialised terminology, which is often perceived as a "new language." Using practice-oriented classroom scenarios, the article illustrates how learners analyse conceptually demanding passages, develop targeted prompts for AI systems, and critically evaluate the generated output. This approach supports both reflection on disciplinary language structures and metacognitive work, and strengthens learners' ability to engage in scientifically grounded thinking. Research findings underline the need for a didactically reflected integration of AI in order to ensure disciplinary accuracy and communicative competence in scientific contexts. Key challenges include safeguarding content accuracy and raising awareness of risks such as declining learner agency and increasing automation effects. The article discusses the opportunities and risks of AI-supported disciplinary language learning and offers practice-oriented suggestions for responsible use that sustainably strengthens scientific communication in a media-literate and collaborative way.

Keywords Didaktische Szenarien, Fachkommunikation, KI-gestütztes Fachsprachenlernen, Kollaboration, Metakognition, Personalisierung, Schreib- und Lesekompetenz

1 Einleitung – KI in der Fachkommunikation

1.1 Aktuelle Entwicklungen und Relevanz Künstlicher Intelligenz für die Fachsprachenvermittlung

Die Integration von Künstlicher Intelligenz (KI) in die Fachkommunikation und Fachsprachenvermittlung verändert die Bildungslandschaft grundlegend. Generative KI unterstützt sprachliche und kommunikative Prozesse im Unterricht und kann diese teilweise automatisieren. Sie übernimmt dabei unterschiedliche Rollen: als Generator für Ideen und Texte, als Transformator für die Vereinfachung oder Übersetzung von Inhalten, als Evaluator für Korrektur und Fehleranalyse sowie als dialogischer Partner (Steinhoff 2023: 12; Tekin 2023: 138). Forschung zeigt, dass KI nicht nur metasprachliche Fähigkeiten stärkt, sondern auch praxisorientierte Lernpläne unterstützen kann, etwa durch die automatisierte Gliederung von Inhalten (Cieżka 2024: 386f.). Besonders im Schreiben und Leseverstehen eröffnen sich Potenziale: KI liefert individuelle Schreibimpulse, erstellt Texte auf unterschiedlichen Niveaus, generiert Aufgaben zum Textverständnis und gibt personalisiertes Feedback, wodurch interaktive Lernarrangements entstehen (Steinhoff 2023: 10f.; Göktas 2024: 19).

1.2 Die Bedeutung der Personalisierung und Kollaboration im Zeitalter von KI

Ein zentrales Merkmal des KI-Einsatzes ist die Möglichkeit zur personalisierten Förderung und zur Unterstützung kollaborativer Lernprozesse. KI passt Inhalte dynamisch an individuelle Bedürfnisse an (Göktas 2024: 18) und fördert kreatives, problemlösendes Denken, indem sie kollaborative Szenarien wie gemeinsames Schreiben oder projektorientiertes Arbeiten unterstützt (Martin/Pengel 2024: 2). Diese Entwicklungen stehen in enger Verbindung mit didaktischen Ansätzen, die Lernende nicht mehr als passive Rezipienten, sondern als aktive Gestalter ihrer Lernprozesse begreifen (Lindauer et al. 2013: 5).

1.3 Ziel und Aufbau des Beitrags

Ziel des Beitrags ist es, aktuelle und potenzielle Einsatzmöglichkeiten von KI in der Fachkommunikation zu beleuchten, mit besonderem Fokus auf Personalisierung, Kollaboration und didaktischen Herausforderungen. Der Beitrag gliedert sich in drei Teile: Zunächst werden theoretische Grundlagen zum KI-Einsatz in der Fachkommunikation vorgestellt. Darauf folgen praxisbezogene Einsatzszenarien. Abschließend werden zentrale Potenziale und Limitationen reflektiert sowie Implikationen für die Weiterentwicklung fachsprachlicher Kompetenzen und offene Forschungsfragen diskutiert.

2 Fachsprachenlernen im Wandel: Herausforderungen und Chancen

Die Aneignung von Fachsprache stellt Lernende, insbesondere in heterogenen Gruppen, vor erhebliche Barrieren. Fachsprache fungiert als „neue Sprache“ (Michalak 2015: 139), deren Muster, Strukturen und Vokabular sich deutlich von der Alltagskommunikation unterscheiden. Der Aufbau dieser *Cognitive Academic Language Proficiency* (CALP), die primär schriftsprachlich orientiert ist, kann bis zu sieben Jahre dauern (Neumann 2017: 11; Schmölzer-Eibinger 2008: 48), was die hohen Anforderungen an sprachliche Kompetenzen im Fachunterricht verdeutlicht (Geist/Krafft 2017: 124).

Der Kompetenzerwerb verläuft von einem rezeptiven zu einem produktiven Umgang mit Fachtexten (Michalak 2015: 139). Zunächst steht das Textverständnis im Vordergrund, etwa durch Zusammenfassen oder das Übertragen in Darstellungsformen wie *Mindmaps*. Diese helfen beim Strukturieren und Erkennen von Schlüsselbegriffen, müssen jedoch durch weiterführende Aufgaben, etwa *Concept Maps*, ergänzt werden, um ein vertieftes Textverständnis zu erreichen (Nückles et al. 2004: 51–53).

Produktives Umgestalten von Texten fördert die Fähigkeit, Inhalte adressatengerecht darzustellen und die eigene Textkompetenz zu erweitern (Schmölzer-Eibinger 2008: 15). Im weiteren Verlauf müssen Lernende eigene Fachtexte

verfassen und sich aktiv an Fachgesprächen beteiligen. Dabei zeigen Studien, dass mündliche Ausdrucksmöglichkeiten nicht automatisch die komplexeren Anforderungen schriftlicher Kompetenz abbilden (Geist/Krafft 2017: 124).

Über sprachliche Fertigkeiten hinaus sind metakognitive und kollaborative Fähigkeiten entscheidend (Michalak 2015: 129). Die Literale Didaktik setzt hier an, indem sie Textkompetenz durch aktives Sprachhandeln, kooperatives Schreiben und Reflexion fördert (Schmölzer-Eibinger 2008: 178–183). Planen, Überarbeiten und Reflektieren gehören ebenso dazu wie der Umgang mit unbekannten Begriffen und die Zusammenarbeit in Gruppen. Schreiben dient dabei nicht nur der Wissensweitergabe, sondern auch der Wissenskonstruktion und Selbstreflexion (Schmölzer-Eibinger/Thürmann 2015: 9; Thürmann et al. 2015: 35f.).

Kollaborative Arbeitsformen wie das gemeinsame Verfassen und Diskutieren von Fachtexten fördern Interaktions- und Argumentationsfähigkeit. Sie entschleunigen den Schreibprozess, regen metakognitive Aktivitäten an und unterstützen die Bewältigung komplexer Aufgaben (Geist/Krafft 2017: 134–136).

Zusammenfassend ist zu konstatieren, dass die erfolgreiche Integration von Fachsprache einen systematischen Ausbau rezeptiver und produktiver Fähigkeiten, idealerweise im Rahmen der Literalen Didaktik erfordert. Diese verbindet Sprach- und Fachlernen, stärkt kooperatives Sprachhandeln und adressiert die Herausforderungen heterogener Lerngruppen mit unterschiedlichen sprachlichen Voraussetzungen.

3 Einsatzfelder von KI im Fachunterricht

Der Einsatz generativer KI-Tools eröffnet im naturwissenschaftlichen Unterricht vielfältige Möglichkeiten zur Unterstützung, Individualisierung und Vertiefung von Lernprozessen (Tekin 2023: 138; Cieżka 2024: 375). KI kann dabei unterschiedliche Rollen einnehmen, die über klassische Lehrmethoden hinausgehen und die Entwicklung wissenschaftlicher Arbeitsweisen, präzi-

ser Fachsprache und fachlich korrekt formulierter Argumentationen fördern (Steinhoff 2023: 10–12).

Als *Generator* liefert KI Impulse für Versuchsbeschreibungen und Entwurfstexte. Sie unterstützt die Strukturierung komplexer naturwissenschaftlicher Inhalte, erleichtert die Formulierung fachlich korrekter Beobachtungen und Ergebnisse und unterstützt die Planungsphase beim Schreiben von Versuchsprotokollen (Steinhoff 2023: 11; Cieżka 2024: 387).

In der Rolle des *Transformers* kann KI Texte zusammenfassen, umformulieren oder vereinfachen, dabei den Gebrauch von Fachbegriffen konsistent halten und wissenschaftliche Textsorten wie Versuchsprotokolle adäquat den Richtlinien anpassen. So werden metasprachliche Kompetenzen geschult, und Lernende entwickeln die Fähigkeit, komplexe naturwissenschaftliche Inhalte präzise und verständlich zu kommunizieren (Tekin 2023: 156; Göktas 2024: 18).

Als *Evaluator* übernimmt KI die Prüfung von Texten auf fachliche Genauigkeit, Struktur und Verständlichkeit. Sie liefert Feedback zu korrekter Verwendung wissenschaftlicher Terminologie, erkennt Redundanzen, schlägt Gliederungen vor und unterstützt die metakognitive Reflexion über das eigene Schreiben (Philipp et al. 2014: 17; Tekin 2023: 138).

In der Rolle des *Kommunikators* agiert KI als Dialogpartner, beantwortet fachbezogene Fragen, simuliert Diskussionen zu naturwissenschaftlichen Problemstellungen und regt Problemlösungsprozesse an. Dies fördert eigenständiges, fachlich fundiertes Denken und stärkt die korrekte Anwendung der Fachsprache (Steinhoff 2023: 12; Göktas 2024: 18). Zudem kann in Verbindung mit kollaborativem Arbeiten direkt auf individuelle Fragestellungen und Bedürfnisse der Gruppe eingegangen werden. Die Kombination dieser Funktionen fördert kognitive Prozesse, entwickelt kreatives Denken und kollaboratives Arbeiten im naturwissenschaftlichen Kontext (Cieżka 2024: 375; Martin/Pengel 2024: 2). Besonders vorteilhaft ist die Option der adaptiven Anpassung von Aufgaben, Texten und Schwierigkeitsgraden an individuelle Lernvoraussetzungen der Schülerinnen und Schüler, wodurch ihre Motivation und ihr selbstgesteuertes Lernen gestärkt werden. Lernende können sich komplexe naturwissenschaftliche Inhalte in korrekter Fachsprache im eigenen Lerntempo erschließen (Göktas 2024: 31–33). Zudem erleichtert die Nutzung

von KI eine Strukturierung naturwissenschaftlicher Texte und fördert so gezielt auch metasprachliche Kompetenzen. Effektives Prompting zur Recherche, Analyse und Strukturierung fachlicher Inhalte trägt dazu bei, die Qualität von Aussagen kritisch zu prüfen und zu hinterfragen (Steinhoff 2023: 10f.). Auf diese Weise können zentrale Aspekte des naturwissenschaftlichen Arbeitens wie systematische Informationsbewertung und präzise Wissenschaftskommunikation zielgruppengerecht entwickelt werden.

4 Didaktische Szenarien und die Integration von KI

Didaktische Szenarien zielen darauf ab, KI-gestützte Aufgabenstellungen systematisch in schulische Kontexte zu integrieren, um Lernprozesse gezielt zu fördern und fachsprachliche Kompetenzen aufzubauen. Komplexe Problemstellungen oder sprachlich anspruchsvolle Textstellen dienen dabei als Ausgangspunkt für vertiefte Lernprozesse. Durch präzise Prompts werden Lernende befähigt, sich sowohl fachlich als auch sprachlich mit Inhalten auseinanderzusetzen und geeignete Ausdrucksformen zu entwickeln. Eine anschließende metakognitive Reflexion über kooperative Lernsettings unterstützt die Integration unterschiedlicher Perspektiven und fördert kompetenzorientiertes Lernen.

KI kann ihr Potenzial entlang des gesamten Schreib- und Leseprozesses entfalten, indem sie technisch als *Ghostwriter* erste Textentwürfe generiert, als *Writing Partner* beim Formulieren und Überarbeiten unterstützt oder als *Evaluator* gezielte Rückmeldungen gibt. Besonders die Überarbeitungsphase ermöglicht durch Korrekturvorschläge und Fehleranalysen eine vertiefte Auseinandersetzung mit sprachlichen Strukturen und unterstützt den Aufbau metasprachlicher Fähigkeiten (vgl. Geist/Krafft 2017: 84f.). Im Fachunterricht wird dieser Einsatz durch Textmuster, Satzbaukästen oder spezifische Formulierungshilfen ergänzt, wodurch eine Progression von gelenkter zu freier Textproduktion gezielt gefördert wird, woraus sich schließlich die notwendige Schreibkompetenz entwickelt.

4.1 Individuelle Förderung des Leseverständnisses von Fachtexten

Die Förderung des Leseverständnisses fachsprachlicher Texte ist eine zentrale Herausforderung im Fachunterricht. KI kann durch personalisierte Texte auf unterschiedlichen Schwierigkeitsstufen unterstützend wirken, indem sie Sprachniveau, Textsorte, Länge oder thematische Schwerpunkte anpasst. Lernende erhalten so die Möglichkeit, komplexe Fachtexte schrittweise zu erschließen und sprachliche Hürden zu überwinden.

Darüber hinaus können KI-Tools gezielt Aufgaben wie W-Fragen, Lückentexte oder Multiple-Choice-Aufgaben generieren, die nicht nur das inhaltliche, sondern auch das sprachlich-strukturelle Verständnis fördern. Automatisiertes Feedback ermöglicht dabei ein eigenständiges Arbeiten im individuellen Lerntempo. Dies trägt insbesondere zur fachsprachlichen Lesemotivation bei, da Lernende erleben, dass sie komplexe Inhalte zunehmend sprachlich kontrollieren und verstehen können (Tekin 2023: 156).

4.2 Kollaboratives Arbeiten und Diskursförderung

Kollaborative Arbeitsformen sind für die Entwicklung fachsprachlicher Diskurskompetenzen relevant. Die gemeinsame Textproduktion, etwa durch KI-gestützte Transformation von Texten in andere Formate oder Schreibstile, eröffnet vielfältige Gelegenheiten zur sprachlichen Aushandlung von Bedeutungen. Lernende diskutieren Formulierungsalternativen, reflektieren die fachliche Präzision und entwickeln ein Bewusstsein für Textsortenkonventionen.

Dieser diskursive Prozess verlangsamt bewusst den Schreibprozess und ermöglicht metakognitive Reflexion über (Fach-)Sprache, Ausdrucksweisen und Argumentationsstrukturen. KI unterstützt dabei als Impulsgeber, indem sie alternative Formulierungen anbietet, die von den Lernenden kritisch bewertet und sprachlich adaptiert werden. So entstehen Lerngelegenheiten, in denen Sprache als Werkzeug des Wissensaufbaus erfahren wird (vgl. Cieżka 2024: 386f.).

4.3 Förderung von Kommunikations-, Schreib- und Lesekompetenz

Der Einsatz von KI fördert die systematische Verknüpfung von Kommunikations-, Schreib- und Lesekompetenzen, die im Fachunterricht nicht isoliert, sondern als wechselseitig stützende Prozesse betrachtet werden. KI kann Schreibansätze generieren, Texte sprachlich anpassen oder Lesehilfen bereitstellen und so den Zugang zu komplexen Fachinhalten erleichtern. Durch personalisierte Impulse und adaptive Lernangebote arbeiten Lernende zunehmend eigenständig und reflektiert. Besonders die Auseinandersetzung mit fachsprachlicher Präzision – etwa beim Erstellen von *Mindmaps*, *Concept Maps* oder beim Umschreiben von Text(sort)en – stärkt metasprachliches Denken und die Fähigkeit, Fachinhalte adressatengerecht zu kommunizieren. Reflexive Schreibprozesse, die von KI unterstützt und durch Peer-Feedback ergänzt werden, ermöglichen zudem den Aufbau eines Repertoires an fachsprachlichen Ausdrucksmitteln (vgl. Geist/Krafft 2017: 85).

4.4 Medienkompetente Lernkultur

Die Integration von KI in den Fachunterricht fördert nicht nur fachliche und sprachliche Kompetenzen, sondern auch eine medienkompetente Lernkultur. Lernende erwerben die Fähigkeit, digitale Werkzeuge bewusst und kritisch für fachsprachliche Lernprozesse einzusetzen. Durch den Einsatz von KI-Tools in authentischen Kommunikationssituationen entwickeln sie ein Bewusstsein für Chancen und Grenzen der Technologien, reflektieren fachsprachliche Inhalte kritisch und übernehmen Verantwortung im Umgang mit KI-basierten Lernangeboten. So wird der fachliche Wissenserwerb mit der Kompetenz verbunden, Sprache, Medien und Wissen in einer digitalisierten Gesellschaft sicher zu nutzen (Tekin 2023: 138; Göktas 2024: 31–33).

4.5 Didaktisches Szenario: KI-gestütztes Erarbeiten und Zusammenfassen eines Fachtexts im Biologieunterricht

Ein Beispiel soll den Mehrwert von KI für fachsprachliches Lernen verdeutlichen: Die Schülerinnen und Schüler erhalten einen Fachtext zur Photosynthese, den sie in einer Zusammenfassung verständlich in eigenen Worten formulieren sollen. Hierfür markieren sie zentrale oder unklare Begriffe und nutzen gezielte Prompts, um von der KI vereinfachte Erklärungen, Definitionen oder Paraphrasen zu erhalten. So werden fachliche Konzepte sprachlich nachvollzogen und adressatengerecht aufbereitet (vgl. Geist/Krafft 2017: 84f.). Dies unterstützt auch die bereits erwähnte Ausbildung von CALP-Kompetenzen.

Darüber hinaus liefert die KI den Lernenden alternative Darstellungen und Beispielsätze, die die Lernenden kritisch mit dem Original vergleichen oder übernehmen können. Dies unterstützt im Sinne der Literalen Didaktik das Verständnis naturwissenschaftlicher Sprachmuster (Tekin 2023: 156). Anschließend verfassen sie eine eigene präzise Zusammenfassung, reflektieren den Nutzen der KI-Hilfen und trainieren so den korrekten, adressatengerechten Sprachgebrauch (Ciężka 2024: 386f.). An dieser Stelle kann die KI als Evaluator eingesetzt werden, um die Texte nochmals zu überarbeiten und damit sowohl CALP als auch Textverständnis nochmals zu forcieren. Das Szenario stärkt Wortschatzverstehen (CALP), Textverständnis, Ausdrucksfähigkeit und kritisches Denken und zeigt, wie Fachsprache im naturwissenschaftlichen Unterricht aktiv erarbeitet und angewendet werden kann.

5 Herausforderungen durch die Verwendung von KI im Fachunterricht

Die Integration von KI in den Fachunterricht eröffnet vielfältige Chancen, bringt jedoch auch spezifische Herausforderungen mit sich, die didaktisch reflektiert adressiert werden müssen. Ein zentrales Problem liegt in der Sicherung von fachlicher Präzision und sprachlicher Genauigkeit. KI-generierte

Texte müssen kontinuierlich überprüft werden, da ungenaue Formulierungen oder fachlich falsche Inhalte übernommen werden könnten. Eine systematische Qualitätskontrolle ist daher unverzichtbar, um die fachsprachliche Korrektheit zu gewährleisten (Göktas 2024: 31–33).

Darüber hinaus birgt die Automatisierung von Schreibprozessen die Gefahr, dass Lernende Aufgaben zunehmend an die Technologie delegieren und dadurch an Eigenaktivität verlieren. Dies kann die Entwicklung fachsprachlicher Ausdruckskompetenz hemmen, wenn Schülerinnen und Schüler Texte lediglich übernehmen, ohne diese kritisch zu reflektieren oder sprachlich weiterzuentwickeln. Um dieser Tendenz entgegenzuwirken, ist es notwendig, KI-Aufgaben so zu gestalten, dass aktive Sprachproduktion und Reflexion zwingend erforderlich bleiben.

Im Kontext der fachsprachlichen Förderung durch KI ist besonders bedeutsam die gezielte Förderung von Prompt-Kompetenz. Sie umfasst die Fähigkeit, Aufgabenstellungen in präzise, fachsprachlich korrekte und zielführende Eingaben für KI-Systeme zu übersetzen. Dazu gehört nicht nur die bewusste Verwendung von Fachterminologie und klaren Formulierungen, sondern auch die Fähigkeit zur Reflexion, Bewertung und Anpassung der Eingaben. Didaktisch lässt sich diese Kompetenz schrittweise aufbauen, indem Lehrkräfte zunächst modellhaft die Wirkung unterschiedlicher Prompts verdeutlichen, bevor Lernende eigene Eingaben formulieren, kritisch prüfen und im Austausch mit anderen weiterentwickeln. In der eigenständigen Anwendung reflektieren sie schließlich, welche Formulierungen zu fachlich präzisen Ergebnissen führten. Auf diese Weise wird KI nicht zum Ersatz für Sprachproduktion, sondern zum Werkzeug einer aktiven Auseinandersetzung mit Fachtexten.

Zusammenfassend zeigt sich, dass der Einsatz von KI im Fachunterricht nur dann sein volles Potenzial entfalten kann, wenn er kontinuierlich reflektiert und kritisch begleitet wird. Eine didaktisch durchdachte Einbettung ist erforderlich, um Risiken wie mangelnde Eigenaktivität, ungenaue Fachsprache oder unpräzise Prompting-Praktiken zu vermeiden und die Chancen für die Förderung von Lese-, Schreib- und Fachkompetenzen wirksam nutzbar zu machen.

6 Ausblick

Der Einsatz von Künstlicher Intelligenz (KI) im Fachunterricht eröffnet vielfältige Perspektiven für eine lernförderliche, kollaborative und inklusive Gestaltung von Lehr- und Lernprozessen. KI kann unterschiedliche Rollen einnehmen: Als *Generator* liefert sie Impulse für Ideenfindung oder erstellt personalisierte Texte und Übungen; als *Transformator* ermöglicht sie die Umwandlung von Texten in andere Formate, Schwierigkeitsstufen oder Fachsprachen; als *Evaluator* unterstützt sie durch Korrektur und Feedback; und als *Kommunikator* steht sie beratend zur Seite (Tekin 2023: 138). Besonders als *Writing Tutor* und *Partner* kann KI die Entwicklung von Lese- und Schreibkompetenzen fördern, indem sie differenzierte, an individuelle Lernstände angepasste Aufgaben bereitstellt und kollaboratives Arbeiten unterstützt (Göktas 2024: 18; Martin/Pengel 2024: 2; Steinhoff 2024: 10–12). Dies umfasst sowohl Schreibaufgaben wie Einstiegs- oder Rekapitulationsimpulse als auch Leseverständnisaufgaben, die Lernende im eigenen Tempo bearbeiten können.

Ein besonders wichtiger Aspekt ist die Fachsprache im naturwissenschaftlichen Unterricht. KI kann komplexe Fachtexte aufbereiten, schwierige Terminologie erklären und Lernenden helfen, fachsprachliche Präzision zu entwickeln. Durch gezielte Prompts lassen sich Definitionen, Prozessbeschreibungen oder Fachbegriffe in verständlicher Form aufbereiten, wodurch Lernende diese sicher anwenden und in eigene Texte übertragen können. Die Reflexion über sprachliche Strukturen und Fachtermini wird so systematisch unterstützt, was den Erwerb fachsprachlicher Kompetenz und das Verstehen komplexer naturwissenschaftlicher Zusammenhänge stärkt.

Für die zukünftige Gestaltung des KI-Einsatzes im Unterricht ergeben sich zentrale Anforderungen. In der Forschung gilt es zu klären, wie KI gezielt für die Entwicklung fachsprachlicher Aufgaben, zur Unterstützung des Textverstehens und zur iterativen Förderung von Kompetenzen eingesetzt werden kann (Göktas 2024: 18).

In der Lehrkräfteausbildung sind methodische, technologische und sprachdidaktische Kompetenzen im Umgang mit KI unabdingbar, um Lernprozesse differenziert, adaptiv und verantwortungsvoll zu gestalten. In der Unterrichtspraxis wiederum sollten Szenarien entwickelt werden, die Raum für Kollabo-

ration, Reflexion und individualisiertes Lernen eröffnen und KI gezielt zur Förderung fachsprachlicher Kompetenzen im naturwissenschaftlichen Kontext einsetzen (Göktas 2024: 18; Martin/Pengel 2024: 2).

Nachhaltige Fachkommunikation im Zeitalter von KI erfordert, dass alle Prozessbeteiligten – Forschende, Lehrkräfte und Lernende – gemeinsam an einem reflektierten und verantwortungsvollen Umgang mit KI arbeiten. Dies schließt die Beherrschung präziser Prompt-Formulierungen ebenso ein wie die Förderung von metakognitiven, kooperativen und fachsprachlichen Kompetenzen, die für die kritische Nutzung von KI-Tools unerlässlich sind (Tekin 2023: 138; Göktas 2024: 18–19; Martin/Pengel 2024: 2). Entscheidend wird sein, die Potenziale von KI zur Förderung von Kommunikation, Kreativität und Problemlösefähigkeit mit einer systematischen Sprachreflexion zu verbinden.

Damit Bildung im naturwissenschaftlichen Unterricht im Kontext von KI lernförderlich und zukunftsfähig gestaltet werden kann, muss die Fachsprache konsequent in den Mittelpunkt gestellt werden. KI sollte nicht als Ersatz, sondern als didaktisch gesteuertes Werkzeug verstanden werden. Sie strukturiert Lernprozesse, fördert Sprachbewusstsein für Fachtermini und den sicheren Umgang damit in naturwissenschaftlichen Texten. Durch personalisierte Textaufbereitung, gezielte Prompts und Feedback können Lernende ihre fachsprachliche Präzision, Lesekompetenz und Schreibfähigkeit systematisch ausbauen. So wird KI zu einem zentralen Instrument, um Fachwissen, Sprachkompetenz und metakognitive Fähigkeiten gleichzeitig zu fördern.

Literatur

- CIEŻKA, A. (2024): „Generative KI-Tools: Die Zukunft des kreativen Lernens – Potenzial für kritisches Denken und Kreativität.“ *ZIF – Zeitschrift für Interkulturellen Fremdsprachenunterricht. MehrSprachen Lernen und Lehren* 29.1: 375–405. DOI: 10.48694/zif.3722.
- FACHLEHRPLÄNE REALSCHULE (2017): *Deutsch 5. Gültig ab Schuljahr 2017/18*. Online unter: <https://www.lehrplanplus.bayern.de/fachprofil/realschule/deutsch/5> (20.08.2025).

- GEIST, B. / KRAFFT, A. (2017): *Deutsch als Zweitsprache*. Tübingen: Narr Francke Attempto.
- GÖKTAS, S.A. (2024): „Zur Vermittlung des Leseverstehens im DaF-Unterricht mit KI-Tools am Beispiel von ChatGPT.“ *Trakya Germanistik Studien 3: Künstliche Intelligenz in der Sprache und Bildung*, 15–35.
- GRIESSHABER, W. (2017): „Schreiben in der Zweitsprache Deutsch.“ *Deutsch als Zweitsprache*. (Deutschunterricht in Theorie und Praxis 9). Hrsg. B. Ahrenholz / I. Oomen-Welke. Baltmannsweiler: Schneider Hohengehren. 292–305.
- KONRAD, K. (2009): *Leseförderung in der Sekundarstufe I. Über den Wissenserwerb beim Lesen*. (Studienreihe: Psychologische Forschungsergebnisse 140). Hamburg: Dr. Kovač.
- LINDAUER, T. / SCHMELLENTIN[-BRITZ], C. / BEERENWINKEL, A. / HEFTI CHRIST, C. / FURGER, J. (2013): *Sprachbewusst unterrichten. Eine Unterrichtshilfe für den Fachunterricht*. Basel: Pädagogische Hochschule FHNW Nordwestschweiz.
- MARTIN, A. / PENGEL, N. (2024): „Kreative Nutzung von KI zur Personalisierung des Lernens – Personalisierung & kollaboratives, kreatives Lernen.“ *LBzM – Ludwigsburger Beiträge zur Medienpädagogik* 24: 1–11.
- MICHALAK, M. (2015): *Sprache im Fachunterricht. Einführung in DaZ und sprachbewussten Unterrichten*. Tübingen: Narr Studienbücher.
- NEUMANN, A. (2017): „Theorie und Forschungsstand. Sprach- und Fachlernen im Fach Deutsch und darüber hinaus.“ „!!Fach-an-Sprache-an-Fach!!“ *Schreibförderung in Deutsch und Mathematik*. Hrsg. S. Rossag / A. Neumann / D. Leiss / K. Schwippert. Baltmannsweiler: Schneider Verlag Hohengehren. 9–29.
- NÜCKLES, M. / GURLITT, J. / PABST, T. / RENKL, A. (2004): *Mind Maps & Concept Maps. Visualisieren, Organisieren, Kommunizieren*. München: dtv.
- PHILIPP, M. / BRÄNDLI, M. / KIRCHHOFER, K. C. (2014): *Kooperatives Lesen. Lesefluss, Textverstehen und Lesestrategien verbessern*. Seelze: Klett-Kallmayer.
- SCHMÖLZER-EIBINGER, S. (2008): *Lernen in der Zweitsprache. Grundlagen und Verfahren der Förderung von Textkompetenz in mehrsprachigen Klassen*. Tübingen: Narr.
- SCHMÖLZER-EIBINGER, S. / THÜRMANN, E. (2015): „Schreiben als Medium des Lernens.“ *Schreiben als Medium des Lernens* (Gesellschaft für Fachdidaktik 8). Hrsg. S. Schmölzer-Eibinger / E. Thürmann. Münster: Waxmann. 9–15.
- STEINHOFF, T. (2023): „Künstliche Intelligenz als Ghostwriter, Writing Tutor und Writing Partner. Zur Modellierung und Förderung von Schreibkompetenzen im

- Zeichen der Automatisierung und Hybridisierung der Kommunikation am Beispiel des Schreibens mit ChatGPT in der 8. Klasse.“ *Personale und funktionale Bildung im Deutschunterricht. Theoretische, empirische und praxisbezogene Perspektiven*. Hrsg. C. Albrecht et al. Stuttgart: Metzler. 85–99. DOI: 10.1007/978-3-662-69640-8.
- TEKİN, Ö. (2023): „ChatGPT im Unterricht Deutsch als Fremdsprache.“ *Alman Dili ve Kültürü Araştırmaları Dergisi* 5.2: 135–163. DOI: 10.55143/alkad.1390420.
- THÜRMANN, E. / PERTZEL, E. / SCHÜTTE, A. U. (2015): „Der schlafende Riese: Versuch eines Weckrufs zum Schreiben im Fachunterricht.“ *Schreiben als Medium des Lernens* (Gesellschaft für Fachdidaktik 8). Hrsg. S. Schmölzer-Eibinger / E. Thürmann. Münster: Waxmann. 17–46.
- TSCHEKAN, K. (2012): *Kompetenzorientiert unterrichten. Eine Didaktik*. Berlin: Cornelsen.

Präsentationen im Zeitalter Künstlicher Intelligenz

Einsatz, Reflexion und Beurteilung im fremd- und fachsprachlichen Unterricht an Hochschulen

Abstract

The genre “presentation” represents a central communicative act in academic settings, particularly in foreign-language and subject-specific language courses accompanying academic studies, where technical content must be appropriately adapted to the target audience in oral format. However, it is precisely in this context that a special focus is needed on the influence and potential of AI in the creation and implementation of presentations. The objective of this article is to discuss how AI-tools — such as those for text generation, visualization, or language analysis — can support the creation, delivery, and reflection of presentations, but also what negative effects they may have on the development of students’ presentation skills in specialized and foreign language university courses. An overview of relevant AI tools that can also be applied in later professional presentation practice and is intended to provide teachers and learners with a practical repertoire of digital aids. The findings provide impetus for the responsible and skill-oriented integration of AI into foreign language teaching.

Keywords fachsprachlicher Unterricht, KI-Tools, Präsentation, studienbegleitender Unterricht

1 Einleitung

Die Textsorte „Präsentation“ stellt ein zentrales kommunikatives Format in akademischen Lernsettings, speziell im fremd- und fachsprachlichen, studienbegleitenden Unterricht dar, in denen fachliche Inhalte adressatengerecht aufbereitet und mündlich vermittelt werden müssen. Gerade in diesem Kontext bedarf es eines besonderen Fokus auf den Einfluss und die Potenziale von Künstlicher Intelligenz (KI) bei der Erstellung und Durchführung.

Ziel des vorliegenden Beitrags ist es, diskursiv darzustellen, wie KI-gestützte Werkzeuge – etwa zur Textgenerierung, Visualisierung oder Sprachanalyse – die Erstellung, Durchführung und Reflexion von Präsentationen unterstützen können, aber auch, welche Nachteile sie für die Entwicklung von Präsentationskompetenzen Studierender im fach- und fremdsprachlichen hochschulischen Unterricht haben. Im Folgenden sollen typische Merkmale der Textsorte anhand einer qualitativen Inhaltsanalyse erfasst und der Einfluss von KI auf die Gestaltung untersucht werden. Zudem werden didaktische Überlegungen zur Gestaltung von Unterrichtsszenarien in Hinblick auf die Integration der Textsorte „Präsentation“ angestellt. Eine Übersicht über relevante und für die spätere Berufspraxis nutzbare KI-Tools, speziell für Präsentationen, soll Lehrkräften wie Lernenden ein praxisorientiertes Repertoire an digitalen Hilfsmitteln bieten. Die Erkenntnisse liefern Impulse für eine verantwortungsbewusste und kompetenzorientierte Integration von KI im fremdsprachlichen Fachunterricht.

2 Die Bedeutung von Präsentationen im akademischen Kontext

Die Präsentation ist eine zentrale akademische Textsorte, deren Ausbildung im allgemeinsprachlichen und fachbezogenen Sprachunterricht an Hochschulen auf verschiedenen Niveaustufen erfolgt. Präsentationen fördern nicht nur die Fähigkeit, Inhalte strukturiert und verständlich aufzubereiten, sondern auch das souveräne Auftreten vor Publikum. Zudem dienen sie als Evaluations- und

Leistungsnachweis, in dem kommunikative und fachliche Kompetenzen sichtbar werden. Damit tragen Präsentationen wesentlich zur wissenschaftlichen Ausbildung und Persönlichkeitsentwicklung bei. Rhein/Henning (2024: 523) bezeichnen „Präsentationen als komplexe Kommunikationsformen, die für Wissenschaftler*innen und Studierende eine Herausforderung darstellen“. Daraus ergibt sich die Frage, warum gerade diese Kommunikationsform speziell im studienbegleitenden Deutsch-als-Fremdsprache- und fremdsprachlichen Unterricht zur Herausforderung für viele Studierende an Hochschulen wird.

Die zunehmende Verbreitung von KI eröffnet für viele Studierende neue Möglichkeiten, sich mit KI-gestützten Werkzeugen Hilfe bei der Erstellung und Durchführung von Präsentationen zu suchen, sei es im Bereich der Textgenerierung, als Übersetzungshilfe oder zur Visualisierung. So gaben im Rahmen einer Befragung des „Monitors Digitalisierung 360°“ vom Wintersemester 2023/2024 bereits knapp die Hälfte der Studierenden an, KI-generative Tools im Studium, u. a. für die Vorbereitung von Präsentationen, zu nutzen (Budde et al. 2024: 64). Diese Technologien können Studierende bei der Strukturierung, sprachlichen Gestaltung und medialen Aufbereitung ihrer Präsentationen unterstützen, werfen jedoch zugleich Fragen zu Urheberschaft, Transparenz, wissenschaftlicher Redlichkeit und Präsentationskompetenz auf.

Um herauszufinden, welche Folgen das verstärkte Zurückgreifen auf KI-gestützte Werkzeuge zur Erstellung dieser Textsorte und dadurch für die Kompetenzgewinnung der Studierenden hat, müssen einerseits die für den Fremdsprachenunterricht relevanten Präsentationskompetenzen identifiziert, andererseits die Relevanz von KI-gesteuerten Werkzeugen zur Erstellung von Präsentationen für Studierende untersucht und auf ihre Vor- und Nachteile hin betrachtet werden. Denn eines der Lernziele im studienbegleitenden Fremd- und Fachsprachenunterricht ist nun mal die Entwicklung der Kommunikationsfähigkeit in der Fremd- und Fachsprache. Entsprechend kann das Präsentieren als „spezielle Aufgabe und besondere kommunikative Interaktion“ gesehen werden, die das „mediengestützte, primär monologische, adressatenorientierte Sprechen“ fördert (Ruth-Herbein et al. 2021: 388).

Vor diesem Hintergrund zielt der vorliegende Artikel darauf ab, die Textsorte „Präsentation“ im akademischen Fremd- und Fachsprachenunterricht unter dem Einfluss KI-gestützter Werkzeuge zu untersuchen. Es soll heraus-

gefunden werden, welche Potenziale und Herausforderungen der Einsatz von KI für die Erstellung, Durchführung und Reflexion von Präsentationen birgt, welche Kompetenzen dabei gefördert oder möglicherweise vernachlässigt werden. Ebenso sollen die Vor- und Nachteile der gängigsten KI-gesteuerten Werkzeuge in diesem Zusammenhang unter Anbetracht einer für die Studierenden kompetenten, reflektierten und wissenschaftlich redlichen Nutzung gegenübergestellt werden. Damit soll ein Beitrag geleistet werden, die Integration von KI im Unterricht nicht nur technisch, sondern auch kompetenzorientiert und kritisch-reflexiv bei der Erstellung und Durchführung von Präsentationen zu gestalten.

3 Theoretische und didaktische Grundlagen

Um festhalten zu können, was Präsentationskompetenz genau ist und welche Auswirkungen KI-gestützte Präsentationen im Unterricht darauf haben, müssen wir uns zuallererst mit der Frage beschäftigen, was konkret eine Präsentation im Hochschulkontext ist. Die Textsorte kann als Herantasten an die Diskursform der Präsentation als wissenschaftliches Informationswerkzeug gesehen werden (Grabowski 2003). In diesem Sinne wäre vor allem das Inhaltliche zu betrachten. Denn die Präsentation im wissenschaftlichen Kontext orientiert sich daran, den „Mitgliedern der eigenen [wissenschaftlichen] Disziplin“ (Rhein/Henning 2024: 523) neue Erkenntnisse und Entwicklungen zu vermitteln. Sieht man sich die Textsorte Präsentation im Fremdsprachenunterricht an, so geht es hierbei nicht primär um die Vermittlung von Wissen und Informationen aus dem eigenen Forschungsbereich. Die Präsentation ist hier auch ein Werkzeug, um die eigenen mündlichen und schriftlichen Ausdrucksfähigkeiten unter Berücksichtigung der angesprochenen Zielgruppe zu verbessern und den aktuellen Stand aufzuzeigen. Auf der Multimodalität als zu Grunde liegender Eigenschaft von Präsentationen (Dyngowska et al. 2012) liegt hier das Hauptaugenmerk. Einerseits werden Präsentationen oft als Textsorte gesehen, die im Fremdsprachenunterricht gerne als Prüfungsformat für den mündlichen Ausdruck verwendet wird. Andererseits ist aber bereits

die rein schriftliche Vorbereitung der Präsentation Teil der Aufgabe, sodass sich hierdurch schon die bei der Durchführung „simultane Ko-Präsenz unterschiedlicher kommunikativer Elemente wie gesprochener und geschriebener Sprache“ (ebd.: 40) zeigt. Angelehnt an Bucher (2012) und Ruth-Herbein et al. (2021: 388) sind die Charakteristika einer wissenschaftlichen Präsentation daher „multimodal, multikodal und multimedial“ geprägt. Die logische und nachvollziehbare Verknüpfung des Schriftlichen und Mündlichen und die Fähigkeit, durch diese Verknüpfung den wissenschaftlichen Inhalt und die wichtigsten Informationen in einer „kohäsions- und kohärenzstiftenden Weise“ (Ruth-Herbein et al. 2021: 389) den Zuhörenden zu vermitteln, scheinen daher die eigentliche Aufgabe des Präsentierenden zu sein. Münzt man diese Annahme auf den Fremdsprachenunterricht um, so kann dadurch auch die schriftliche Kompetenz, genauso wie die mündliche Ausdrucksfähigkeit in der Fremdsprache, durch die Lehrperson abgelesen werden. Diese Kompetenzen gilt es im Fremdsprachenunterricht nicht nur zu entwickeln, sondern auch anzuwenden. Entsprechend kann davon ausgegangen werden, dass Studierende, wenn sie Präsentationen im Rahmen ihres Unterrichts erstellen, die Fähigkeit haben, basierend auf ihrem Sprachniveau, Kohärenz und Kohäsion in der Fremdsprache herzustellen und diese in einer adressatengerechten Form zu präsentieren. Dabei wird deutlich, dass Präsentationen nicht nur als Leistungsnachweis dienen, sondern auch als didaktisches Mittel zur Förderung der kommunikativen Kompetenz in der Fremdsprache. Folglich stellen sie eine Schnittstelle zwischen sprachlicher Korrektheit, inhaltlicher Strukturierung und situationsangemessener Ausdrucksfähigkeit dar, die im Lernprozess gezielt geschult und reflektiert werden sollte (Ruth-Herbein et al. 2021). Das Erzielen von Präsentationskompetenz steht also bei der Anwendung der Textsorte „Präsentation“ im Fremdsprachenunterricht im Vordergrund.

Wie kann man Präsentationskompetenz definieren? Ist sie ausschließlich in der klaren und transparenten, dem Publikum verständlichen multimodalen Verknüpfung von schriftlichen und mündlichen Zeichen zu sehen?

In der Linguistik wird die Präsentationskompetenz als Teil der kommunikativen Kompetenz verstanden, die die Fähigkeit umfasst, sprachliche, inhaltliche und mediale Mittel so zu verknüpfen, dass Informationen adressaten- und situationsgerecht vermittelt werden. Dabei geht es sowohl um sprachliche

Korrektheit als auch um die gezielte Herstellung von Kohärenz und Kohäsion innerhalb mündlicher und schriftlicher Darstellungsformen (Dyngowska et al. 2012). Präsentationskompetenz kann daher als eine Fähigkeit beschrieben werden, Sprache funktional einzusetzen, um wissenschaftliche oder fachliche Inhalte klar, überzeugend und interaktiv zu vermitteln (Ruth-Herbein et al. 2021).

In der Psychologie wiederum kann Präsentationskompetenz vor allem als „Ressourcenpool verstanden werden, der sich auch aus verschiedenen Eigenschaften einer Person zusammensetzt“ (ebd.: 390). Präsentierende Personen benötigen demnach nicht nur das fachliche Wissen, sondern auch die Fähigkeit, von ihrer Kompetenz in ihrem Fachgebiet zu überzeugen und gleichzeitig über ein ausreichendes „Repertoire an Verhaltenspotenzialen“ (ebd.: 391) zu verfügen. Hierunter fallen eben jene Aspekte einer Präsentation, die über die schriftliche und mündliche Ausdrucksfähigkeit hinausgehen, nämlich die nonverbalen Kompetenzen. Sie stellen ebenfalls einen wesentlichen Bestandteil erfolgreicher Präsentationen dar, da sie maßgeblich zur Verständlichkeit, Überzeugungskraft und Wirkung des Vortrags beitragen. Ein selbstsicheres Präsentationsverhalten (u. a. durch eine aufrechte und offene Körperhaltung) stärkt die Glaubwürdigkeit der präsentierenden Person (Burgoon et al. 2022). Auch eine der Situation angepasste Gestik und Mimik kann zur Veranschaulichung komplexer Inhalte beitragen (Hostetter 2011). Schließlich tragen stimmliche Merkmale wie Tonhöhe, Lautstärke, Tempo und Betonung zur prosodischen Gestaltung des Vortrags bei und unterstützen sowohl die inhaltliche Strukturierung als auch die rhetorische Wirkung der Präsentation (Dombrowski/Niebuhr 2005).

Aus Sicht der Fremdsprachendidaktik lässt sich Präsentationskompetenz daher als ein umfassendes Set sprachlicher, medialer und kognitiver Fähigkeiten verstehen, die im weitesten Sinne zu den literalen Kompetenzen zählen. Diese Fertigkeiten sind entscheidend für erfolgreiches gesellschaftliches Handeln, da sie es ermöglichen, komplexe Anforderungssituationen zu bewältigen, die sowohl im hochschulischen als auch in zahlreichen beruflichen Kontexten von erheblicher Relevanz sind (Gätje et al. 2016). Unter dem Aspekt der Vermittlung der Textsorte „Präsentation“ im Fremdsprachenunterricht liegt der Fokus schon lange nicht mehr primär auf der sprachlichen Korrektheit, son-

dern fokussiert sich verstärkt auf die sprachlichen, medialen und kognitiven Fähigkeiten, die zur Entwicklung einer umfassenden Präsentationskompetenz beitragen.

4 KI-gestützte Werkzeuge im Kontext von Präsentationen

Im Zeitalter von KI stellt sich die Frage, ob es durch KI-erstellte Präsentationen, wie zuvor herausgearbeitet, zu einer Vernachlässigung der Förderung von Präsentationskompetenzen kommen kann. Zahlreiche Hochschulen bieten ihren Studierenden Leitfäden zur Erstellung von Präsentationen, sowohl was das Inhaltliche als auch das Formelle und Visuelle betrifft. Eines haben aber alle Leitfäden gemeinsam, nämlich den Hinweis, dass vor jeder Auswahl der Präsentationstools, der inhaltlichen Erarbeitung und der Erstellung des Aufbaus eine Idee stehen muss. Erst dann kann eine inhaltliche Arbeit beginnen. Schon das Präsentationsthema beinhaltet bereits erste Gedanken an die wesentlichen Aspekte einer Präsentation, deren Zahl zwar je nach Leitfaden variieren kann, aber sich generell auf jene Fragen bezieht: „Was soll vorgestellt werden?“, „Was kommt beim Zuhörenden an?“, „Wird dem Gesagten geglaubt?“ und „Passen die Aussagen zu den Daten, die zur Erstellung herangezogen wurden?“

Wenn nun Texte für die Präsentationsfolien sowie die inhaltlichen Strukturen nicht selbst erstellt und währenddessen bereits durchgedacht, bearbeitet und regelmäßig auf ihre Kohäsion und Kohärenz hinterfragt werden, können in der Folge typische Fehler oft darin bestehen, „dass die Vortragenden keine Vorstellung davon haben, was sie eigentlich berichten wollen“ (Förstemann/Löffler 2020: 1). Obwohl die ersten Gedanken zum Präsentationsthema sowie die Strukturierung der Präsentation generell zumeist immer noch durch die Präsentierenden selbst passieren, wird der zweite Schritt – die Textgenerierung – bereits oft mit KI durchgeführt. Große Sprachmodelle wie ChatGPT, Gemini, Claude oder Microsoft Copilot erstellen bei gezieltem Prompting bereits eine auf die Zielgruppe zugeschnittene Struktur sowie Aufbereitung des Themas. Die Prompts können diesbezüglich mittlerweile so exakt geschrieben

werden, dass KI-Werkzeuge sämtliche Phasen des Präsentationsprozesses, gegebenenfalls bereits ab der Ideenfindung, von der visuellen Ausgestaltung bis hin zur Analyse der eigenen Vortragsweise unterstützen können. KI-Werkzeuge bieten Hilfestellung bei der Erarbeitung von Einleitungen und formulieren Argumentationsstrukturen, können aber auch komplette Präsentationstexte verfassen. Auf den ersten Blick lassen sich hier einige Vorteile erkennen, die sich dadurch für den/die Präsentierende/n ergeben, wie Zeitersparnis, Effizienz, professionell wirkende Texte.

Aktuelle Forschungen zeigen, dass präzise formulierte Prompts zunehmend dazu beitragen, KI-gestützte Systeme differenziert zu steuern und damit komplexe, mehrstufige Aufgabenprozesse zu unterstützen. Im Kontext von Präsentationen hat das nicht nur Auswirkungen auf die inhaltliche Ideengenerierung und Strukturierung, sondern auch auf die visuelle Ausgestaltung sowie in weiterer Folge auch auf die Analyse der eigenen Vortragsweise (Federiakin et al. 2024). Je präziser die Anforderung an die KI gestellt wird, desto klarer erfüllen generative KI-Werkzeuge die Erwartungen (vgl. Horn 2023: 120).

Eine Voraussetzung für die effiziente und effektive Verwendung von KI-Werkzeugen und ein gezieltes Prompting zur Unterstützung bei Präsentationen ist jedoch ein ausreichend vorhandenes Fachwissen, welches bei (teil)wissenschaftlichen Präsentationen im Rahmen des Studiums, auch im fremdsprachlichen studienbegleitenden Unterricht, Voraussetzung ist, denn ein zugrundeliegendes Fachwissen über das Präsentationsthema in der Erst- bzw. Studierrsprache ist die Basis für eine inhaltlich schlüssige Präsentation. KI-generierte Inhalte können oft fehlerhaft sein und bei nicht vorhandenem Fachwissen auch nicht sorgfältig überprüft werden. Die Frage, die sich hier stellt, ist, ob die Nutzung von textgenerierenden KI-Systemen zur Erstellung des Aufbaus und der Argumentationsstruktur von Präsentationen die Kohäsions- und Kohärenzkompetenz der Studierenden in der Fremdsprache fördert oder im Gegenteil verringert. Ein entsprechendes Sprachniveau ist hier Voraussetzung. So ist im GeR die Kohärenz erst auf B1-Niveau als Verknüpfung „einer Reihe kurzer, einfacher Einzelelemente zu einer zusammenhängenden linearen Äußerung“ (GeR 2001: Abschnitt „Qualitative Aspekte des mündlichen Sprachgebrauchs, Kohärenz B1“) definiert.

Daher ließe sich nun argumentieren, dass bei Präsentationen auf niedrigerem Niveau (A2) KI-generierte Texte bei Kohäsions- und Kohärenzfragen eine Unterstützung sein können. Nur – können Studierende im fach- bzw. fremdsprachlichen studienbegleitenden Unterricht erkennen, ob die durch KI erstellten Texte tatsächlich in sich schlüssig und logisch verknüpft sind, wenn die Kompetenz der zusammenhängenden linearen Äußerungen (vgl. GeR 2001), teilweise der (Fach)wortschatz nicht bekannt und keine für Präsentationen oft typischen Redemittel und grammatikalischen Sprachstrukturen vorhanden sind? Die generative KI kann zwar durchaus mit gezieltem Prompten Textsorten auf niedrigeren Sprachniveaus erstellen, die für Lernende selbst auf einem höheren A-Niveau verständlich sein können, dennoch lassen sich allein durch das Verstehen von Texten noch keine Kenntnisse zur Kohäsion und Kohärenz weiterentwickeln. „Kohäsive Texte können [zwar] das Sachtextverstehen unterstützen“, Voraussetzung hierfür ist aber eine gute Lesefähigkeit, denn „Rezipienten mit schlechterer Lesefähigkeit und weniger Vorwissen können die Textkohäsion nicht für sich nutzen“ (Schmitz 2015: 47). Hervorgehend aus einer Studie von Tengler/Brandhofer (2025: 25) birgt der „Einsatz von generativer KI [...] das Risiko, dass Studierende sich zu sehr auf automatisierte Inhalte stützen, ohne eigene Analyse- und Synthesefähigkeiten zu entwickeln“.

Gleichzeitig besteht ein zentrales Problem auch darin, dass KI-gestützte Textgenerierung Lernende dazu verleiten kann, Vortragsinhalte nahezu vollständig automatisch erstellen zu lassen, wodurch die gezielte und intensive Auseinandersetzung mit dem Stoff leidet, was wiederum dazu führen kann, dass Präsentationstexte schlechter verinnerlicht werden und so die unmittelbare mündliche Darbietung darunter leidet (Tengler/Brandhofer 2025). Spezifisch auch unter dem Gesichtspunkt, dass, wie bereits besprochen, die Kohäsions- und Kohärenzkompetenz weder gefördert noch genutzt wird, um Texte zu erstellen, und so auch der Inhalt dem Sprachniveau und dem Fachwissen entsprechend reproduziert werden kann. Häufig dadurch entstehende Brüche zwischen geschriebenem und gesprochenem Text durch KI-generierte Formulierungen in Präsentationen, die nicht dem natürlichen Sprachgebrauch und dem aktuellen Sprachstand entsprechen, sodass Probleme bei der Kohärenz und Kohäsion zwischen der schriftlichen Variante und dem Mündlichen in der Umsetzung auftreten können (vgl. Berber Sardinha 2023). So zeigen sich

KI-generierte Texte auch oft als „besonders mit Informationen überladen“ und im Vergleich zu von Menschen erstellten Texten als mit weniger Engagement und Integration (ebd.: 9). Die Vermutung, dass die Verwendung von nicht dem Sprachniveau angepasster Lexik und Sprachstruktur auch zu Diskrepanzen bei von KI generierten Präsentationstexten führen kann, sofern sie nicht auf das entsprechende Sprachniveau heruntergebrochen werden, liegt durchaus nahe und sollte in einer Folgestudie untersucht werden.

Im Kontext von Visualisierung und Präsentationskompetenz kann konstatiert werden, dass Visualisierungen als mediale Stütze sprachlicher Produktion fungieren. Die Entscheidung zwischen standardisierten Präsentationsmustern und einem individuellen stilistischen Zugang verweist dabei auf die Fähigkeit, Medien reflektiert und adressatengerecht einzusetzen. Hier zeigt sich Präsentationskompetenz in der Fremdsprache besonders deutlich, da visuelle Gestaltung und sprachliche Darstellung auf das jeweilige Kompetenzniveau abgestimmt werden müssen, um Verständlichkeit und kommunikative Effektivität zu sichern. Im Kontext von Fremdsprachenunterricht und Hochschuldidaktik kann der Einsatz von KI-Präsentationstools problematisch sein, weil dadurch die entscheidende Fähigkeit zur reflektierten Mediennutzung und adressatengerechten Gestaltung unterlaufen werden kann. Automatische Tools tendieren dazu, standardisierte Layouts und Präsentationsmuster zu verwenden, die selten auf das Sprachniveau und die kommunikativen Bedürfnisse von DaF/DaZ-Lernenden zugeschnitten sind (Visualisieren & Präsentieren 2025).

Der Einsatz standardisiert generierter Präsentationstemplates reduziert die Notwendigkeit, Inhalte sowohl inhaltlich als auch gestalterisch bewusst auszuwählen. Gerade im Kontext des universitären Fremd- und Fachsprachenunterrichts ist es jedoch essenziell, dass sprachliche Formulierungen, Visualisierungen und Layout gemeinsam entwickelt werden, um Sprache, Inhalt und Bild aufeinander abzustimmen und damit Verständlichkeit und kommunikative Effektivität zu sichern, also multimodal kompetent präsentieren zu können (Mohamed et al. 2023).

5 KI-Werkzeuge zur Verbesserung der Präsentationskompetenz

KI-Sprachanalyse-Anwendungen können Formulierungen hinsichtlich Verständlichkeit, Stil, Grammatik und Zielgruppenorientierung überprüfen (vgl. LecEval 2025). Feedbacksysteme wiederum geben Hinweise zu Struktur, Tempo oder rhetorischer Wirkung und können so unter bestimmten Voraussetzungen auch zur Verbesserung der Präsentationsqualität beitragen. KI-Werkzeuge erleichtern die Vorbereitung und bieten vielfältige Hilfen, die bisher oft nur durch individuelle Betreuung erreichbar waren. Während solche Anwendungen bereits heute wertvolle Unterstützung leisten, stellt sich die Frage, wie sie konkret in Präsentationssoftware eingesetzt werden können und welche es zu wählen gilt.

Die klassische PowerPoint-Präsentation ist noch immer das gängigste Tool zur Erstellung von Präsentationen. In PowerPoint selbst gibt es aber aktuell noch keine integrierte und vor allem kostenlose KI-Funktion, was oft dazu führt, dass auf Tools zurückgegriffen wird, die im Internet zur Verfügung stehen und mit denen sich Präsentationen automatisch erstellen lassen, ohne dass eine generative KI, wie ChatGPT, Gemini, Claude oder Microsoft Copilot, in einem vorangegangenen Schritt eingesetzt werden muss (vgl. AutoPresent 2025). Der Einstieg ist bei fast allen Angeboten sehr ähnlich. Zuerst wird ein Thema eingegeben, zu dem eine Präsentation erstellt werden soll. Eine grobe Überlegung dazu, welche Struktur und welche Inhalte enthalten sein sollen, hilft bei der Konkretisierung des Präsentationsvorhabens. Viele Dienste bieten anschließend die Möglichkeit, ein Design oder eine Vorlage auszuwählen. Sobald diese Angabe gemacht wurde, übernimmt die KI den Rest. Sie generiert eine passende Gliederung, sammelt Informationen zu den einzelnen Punkten und filtert die wichtigsten Inhalte heraus. Diese werden dann übersichtlich auf Folien übertragen (vgl. Chip.de 2025; AutoPresent 2025). In weiterer Folge werden einige der gängigsten KI-Werkzeuge zur Erstellung von Präsentationen qualitativ auf die Punkte Textgenerierung, Visualisierung, Personalisierungsmöglichkeit, Quellenangabe und Korrektheit der Daten hin untersucht, um einen kurzen, exemplarischen Überblick über gängige KI-Tools, ohne Anspruch auf Vollständigkeit, zu geben.

Die betrachteten KI-Werkzeuge zur Präsentationserstellung unterscheiden sich in ihrem Fokus auf Gestaltung, Textgenerierung und Personalisierung deutlich.

Decktopus arbeitet vor allem promptbasiert und ermöglicht die schnelle Erstellung strukturierter Präsentationen. Das Tool liefert Gliederungen und grundlegende Textbausteine, bleibt jedoch inhaltlich häufig eher oberflächlich (SlideAudit 2025). Die Gestaltung ist zweckmäßig, aber meist stark templatorientiert, wodurch Präsentationen oft generisch wirken. Eine intensivere Personalisierung ist nur eingeschränkt möglich (Slidepeak 2025).

Presentations.AI funktioniert ähnlich, bietet jedoch etwas flexiblere Layouts (PreGenie 2025). Das Tool eignet sich besonders für Nutzer, die eine gute Mischung von formaler Struktur und optischer Kohärenz benötigen. Dennoch erzeugt auch Presentations.AI oftmals standardisierte Formulierungen, sodass trotz eines guten Grundgerüsts eine inhaltliche Nachbearbeitung passieren muss (Plusai 2025).

Beautiful.ai hingegen legt seinen Schwerpunkt auf visuelle Gestaltung. Durch „Smart Slides“ passen sich Layouts automatisch an, sodass auch komplexere Inhalte strukturiert und visuell hochwertig dargestellt werden (SlideAudit 2025). In Verbindung mit themenspezifischen Logos und Farbschemata eignet es sich besonders gut für fachspezifisch gebrandete Präsentationen und liefert gegenüber Decktopus oder Presentations.AI weniger generische Ergebnisse (Beautiful.ai 2025).

Gamma verfolgt einen anderen Ansatz und optimiert Präsentationen stark für digitales Storytelling (Gamma 3.0 2025). Es generiert aus kurzen Stichworten ganze Narrative und eignet sich insbesondere für interaktive oder webbasierte Präsentationen (PreGenie 2025). Zwar entstehen dadurch möglicherweise aufmerksamkeitsstarke Folien, doch im Vergleich zu anderen Tools fehlt es Gamma an Tiefe, weshalb auch hier Präsentationen inhaltlich generisch wirken können, wenn sie nicht weitergehend personalisiert werden (Gamma 2025).

Microsoft 365 Copilot kann anhand von vorgenerierten Texten, aber auch fertigen Dokumenten nicht nur Folien, sondern auch Sprechertexte generieren (DOC2PPT 2022). Dadurch entsteht im Idealfall kein generischer, sondern kontextadaptierter Output, der Fachwissen einbezieht. Einschränkend wirkt,

dass Copilot in der Regel kostenpflichtig ist und eine Überprüfung der Aussagen zwingend notwendig bleibt, besonders wenn externe Quellen genutzt werden, da KI-Modelle fehlerhafte oder ungesicherte Fakten erzeugen können (Copilot 2025).

Prezi fokussiert besonders auf die gestalterischen Aspekte, da es Präsentationen nicht als Folienfolge, sondern zoombar aufbereitet. Visuelle Verknüpfungen, dynamische Übergänge und nichtlineare Navigation machen Inhalte erlebbarer, allerdings eignet sich Prezi weniger für klassische, textlastige Präsentationen oder umfangreiche Datendarstellungen (PreGenie 2025). Die KI-Textgenerierung fällt im Vergleich zu anderen KI-Werkzeugen geringer aus, wodurch Nutzer Inhalte stärker selbst formulieren müssen (Prezi 2025).

Canva wiederum kombiniert KI-gestützte Texterstellung mit Fokus auf Design. Mit *Magic Write* und *Magic Design* können Präsentationen schnell erstellt und anschließend personalisiert und die Visualisierung durch KI-Bilder und Grafiken unterstützt werden (Janßen 2024). Canva kann sowohl professionell als auch generisch wirken, abhängig davon, wie viel an eigener Anpassung vorgenommen wird (Canva 2025).

Ein wesentlicher Faktor, der bei allen Tools beachtet werden muss, ist die Korrektheitsprüfung der Inhalte. Obwohl viele KI-Systeme bereits strukturierte Informationen generieren können, prüfen sie Aussagen nicht zuverlässig auf fachliche Richtigkeit. Daher bleibt eine sorgfältige Validierung zwingend erforderlich, insbesondere bei Zahlen, wissenschaftlichen Fakten oder Quellenbezügen (DOC2PPT 2022). Wenn die Präsentation auf gesicherten Daten basieren soll, müssen Quellen manuell ergänzt und Inhalte fachlich verifiziert werden. KI-Tools beschleunigen zwar die Erstellung von Präsentationen, ersetzen aber nicht die inhaltliche und wissenschaftliche Verantwortung des Erstellers.

6 Fazit und Ausblick

Zusammenfassend lässt sich sagen, dass trotz eines großen Angebots an KI-Werkzeugen zur Erstellung von Präsentationen das Präsentieren mehr ist als

das bloße Gestalten passender Folien, das Erstellen eines Inhaltsverzeichnisses, Generieren eines Textes und in weiterer Folge das Ablesen eines gut formulierten Textes. Es umfasst den Aufbau einer sprechenden Dramaturgie, das aktive Herstellen von Verständlichkeit, die Anpassung an das Publikum, den gezielten Einsatz nonverbaler Elemente und die Fähigkeit, flexibel auf Rückfragen oder Unsicherheiten zu reagieren. Werden diese Fähigkeiten durch automatisierte Prozesse ersetzt, fehlt Lernenden langfristig ein essenzielles Set an kommunikativen und kognitiven Kompetenzen.

Dennoch bieten KI-Werkzeuge bei sinnvoller, zielgerichteter Nutzung auch Funktionen, aus denen sich Chancen und Potenziale für Lernende und Lehrende ergeben. So erhalten Lernende unter Umständen rasch Unterstützung und unmittelbares Feedback, können Präsentationen effizienter vorbereiten und profitieren von einer professionell wirkenden Aufbereitung ihrer Inhalte. Besonders Visualisierungstools ermöglichen einen Zugang zu Präsentationsdesigns, die auch ohne spezifische Fachkenntnisse benutzt werden können. Lehrkräfte können von der schnellen Entstehung erster Entwürfe profitieren, Präsentationssaufgaben stärker differenzieren und Lernprozesse gezielter begleiten.

In Zukunft wird es von entscheidender Relevanz sein, ein Gleichgewicht zwischen dem Einsatz digitaler Hilfsmittel und der Förderung grundlegender Präsentations- und Kommunikationsfähigkeiten zu etablieren. KI kann dabei als unterstützendes Instrument dienen, sollte jedoch nicht die aktive Auseinandersetzung mit Inhalten und die persönliche Entwicklung kommunikativer Kompetenzen ersetzen.

Langfristig liegt das Potenzial darin, Lernende nicht nur bei der Erstellung von Präsentationen zu unterstützen, sondern ihnen auch die Reflexion über die eigene Wirkung, die Strukturierung von Inhalten und die zielgerichtete Ansprache des Publikums zu erleichtern und so Präsentationskompetenzen aktiv zu fördern.

Abschließend sollte generell die Frage gestellt werden, inwiefern Präsentationen, besonders wenn sie vorbereitend erstellt wurden, in Zeiten von KI noch eine sinnvolle und bewertbare Methode zur Leistungsmessung darstellen. Gerade in einem Unterricht, in dem Präsentationen nicht nur als Leistungsnachweis, sondern auch zur Förderung von Präsentationskompetenz dienen,

gewinnt diese Überlegung an Bedeutung. Es wäre daher zu reflektieren, wie Lernende im Prozess des Präsentationserstellens stärker begleitet werden können, um die eigene Präsentationskompetenz gezielt zu entwickeln. Auch könnte der Fokus verstärkt auf den Prozess und nicht ausschließlich auf das Endprodukt gelegt werden, um sicherzustellen, dass Studierende tatsächlich Präsentationsfähigkeiten erwerben. Darüber hinaus wäre es eine Überlegung wert, herauszufinden, wie spontane, kollaborative oder unterrichtsintegrierte Präsentationsphasen dazu beitragen können, individuelle Fähigkeiten sichtbarer zu machen und den Einfluss von KI auf das Ergebnis zu relativieren.

Literatur

- BERBER SARDINHA, T. (2023): „AI-generated vs Human-authored Texts: A Multidimensional Comparison.“ *Applied Corpus Linguistics* 4.1: 100083. DOI: 10.1016/j.acorp.2023.100083.
- BUCHER, H.-J. (2012): „Multimodalität – ein universelles Merkmal der Medienkommunikation: Zum Verhältnis von Medienangebot und Medienrezeption.“ *Interaktionale Rezeptionsforschung. Theorie und Methode der Blickaufzeichnung in der Medienforschung*. Hrsg. H.-J. Bucher / P. Schumacher. Wiesbaden: Springer. 51–82.
- BUDDE, J. / TOBOR, J. / FRIEDRICH, J.-D. (2024): „Monitor Digitalisierung 360°. 2023/24. Wo stehen die deutschen Hochschulen?“ *Arbeitspapier 83 – Künstliche Intelligenz: Einsatz in Studium und Lehre*. Hrsg. Geschäftsstelle Hochschulforum Digitalisierung beim Stifterverband für die Deutsche Wissenschaft e.V. Berlin.
- BURGOON, J. K. / GUERRERO, L. K. / MANUSOV, V. (2022): *Nonverbal Communication* (2nd ed.). London: Routledge.
- DOMBROWSKI, E. / NIEBUHR, O. (2005): „Acoustic Patterns and Communicative Functions of Phrase-final F₀ rises in German: Activating and Restricting Contours.“ *Phonetica* 62.2–4: 176–195.
- DYNKOWSKA, M. / LOBIN, H. / ERMAKOVA, V. (2012): „Erfolgreich Präsentieren in der Wissenschaft? Empirische Untersuchungen zur kommunikativen und kognitiven Wirkung von Präsentationen.“ *Zeitschrift für Angewandte Linguistik* 57.1: 33–65.
- FEDERIAKIN D. / MOLEROV D. / ZLATKIN-TROITSCHANSKAIA, O. / MAUR, A. (2024): „Prompt Engineering as a New 21st Century Skill.“ *Frontiers in Education* 9:1366434. DOI: 10.3389/educ.2024.1366434.

- FÖRSTEMANN, T. / LÖFFLER, A. (2020): *Leitfaden zur Erstellung von wissenschaftlichen Präsentationen*. Berlin: FU Berlin. <https://www.wiwiss.fu-berlin.de/fachbereich/bwl/pruefungs-steuerlehre/loeffler/Lehrmaterialien/Praesentationsleitfaden.pdf> (23.10.2025).
- GÄTJE, O. / KRELLE, M. / BEHRENS, U. / GRUNDLER, E. (2016): „Präsentieren als literale Kompetenz?“ *Leseforum* 1: 5–20.
- GRABOWSKI, J. (2003): „Kongressvorträge und Medieneinsatz: ein Plädoyer für Professionalität.“ *Zeitschrift für Angewandte Linguistik* 39: 53–73.
- HORN, CH. (2023): „Im Dialog mit ChatGPT – 3 Vier potenzielle Einsatzbereiche des Chatbots im Sprachunterricht und beim Fremdsprachenlernen im Test.“ *DaFSzene Korea* 56: *Performative Methoden*: 109–122.
- HOSTETTER, A. B. (2011): „When do Gestures Communicate? A Meta-analysis.“ *Psychological Bulletin* 137.2: 297–315.
- MOHAMED, A. A. / MOHD RADZUAN, N. R. / FAUZI, W. J. / HARBI, S. (2023): „Anxiety Could be a Good Sign: Understanding Challenges in Developing Engineering Graduates' Technical Oral Presentation Competencies.“ *International Journal of Language Education and Applied Linguistics* 13.1: 69–79. DOI: 10.15282/ijleal.v13i1.7531.
- RHEIN, L. / HENNING, U. (2024): „Die Präsentation im Kontext von Deutsch als Fach- und Fremdsprache.“ *Handbuch Deutsch als Fach- und Fremdsprache: Ein aktuelles Handbuch zeitgenössischer Forschung*. Hrsg. M. Szurawitzki / P. Wolf-Farré. Berlin/ Boston: De Gruyter. 523–538.
- RUTH-HERBEIN, E. / LIPPARDT, C. / RINGEISEN, T. (2021): „Konzeptualisierung und Diagnostik von Präsentationskompetenz. Gruppe. Interaktion. Organisation.“ *Zeitschrift für Angewandte Organisationspsychologie (GIO)* 52.4: 387–403.
- SCHMITZ, A. (2015): „Globale Textkohäsion zur Förderung des Textverständnisses. Wie wirkt die globale Textkohäsion bei unterschiedlichen Lesestrategiekenntnissen?“ *Didaktik Deutsch: Halbjahresschrift für die Didaktik der deutschen Sprache und Literatur* 20 (2015) 39: 42–60. https://www.pedocs.de/volltexte/2019/17141/pdf/Didaktik_Deutsch_2015_39_Schmitz_Globale_Textkohaesion.pdf (09.09.2025).
- TENGLER, K. / BRANDHOFER, G. (2025): „KI oder nicht KI? Generative KI – eine Unterstützung des (wissenschaftlichen) Schreibprozesses?“ *Medienimpulse* 63.1: 1–30. DOI: 10.21243/mi-01-25-11.

Visualisieren & Präsentieren (2025): Informationsportal Hochschullehre: Universität Bremen. <https://www.uni-bremen.de/informationsportal-hochschullehre/lehre-gestalten/durchfuehrung/visualisieren-und-praesentieren> (11.11.2025).

Onlinequellen

AUTOPRESENT (2025): https://openaccess.thecvf.com/content/CVPR2025/papers/Ge_AutoPresent_Designing_Structured_Visuals_from_Scratch_CVPR_2025_paper.pdf (27.11.2025).

BEAUTIFUL.AI (2025): <https://www.beautiful.ai/blog/lock-in-branding-with-the-custom-theme-builder-and-new-collection-colors> (20.11.2025).

CANVA (2025): <https://www.theverge.com/news/706637/canva-anthropic-claude-ai-mcp-support> (27.11.2025).

CHIP.DE (2025): https://praxistipps.chip.de/mit-ki-powerpoint-erstellen-die-besten-kostenlosen-tools_178121 (10.10.2025).

COPILLOT (2025): <https://learn.microsoft.com/de-de/copilot/microsoft-365/microsoft-365-copilot-overview> (23.11.2025).

DECKTOPUS (2025): <https://www.decktopus.com/> (01.10.2025).

DOC2PPT (2022): <https://cdn.aaii.org/ojs/19943/19943-13-23956-1-2-20220628.pdf> (10.11.2025).

GAMMA (2025): <https://gamma.app/de/ai-presentation-generator> (23.11.2025).

GAMMA 3.0 (2025): <https://gamma.app/de/insights/introducing-gamma-3-0> (29.11.2025).

GeR (2001): <https://www.europaeischer-referenzrahmen.de/sprachkenntnisse.php> (01.11.2025).

JANSSEN, J. (2024): Präsentationen erstellen mit Canva Magic Studio. <https://www.joergjanssen.com/blog/pr%C3%A4sentationen-erstellen-mit-canva-magic-studio> (10.11.2025).

LECEVAL (2025): <https://arxiv.org/abs/2505.02078> (27.11.2025).

PREGENIE (2025): <https://aclanthology.org/2025.findings-emnlp.165.pdf> (27.11.2025).

PRESENTATIONS.AI (2025): <https://www.presentations.ai/> (17.12.2025).

PREZI (2025): <https://prezi.com/de/features/ai> (15.11.2025).

PLUSAI (2025): <https://plusai.com/blog/in-depth-review-of-presentations-ai> (20.11.2025).

SLIDEAUDIT (2025): <https://dl.acm.org/doi/10.1145/3746059.3747736> (28.11.2025).

SLIDEPEAK (2025): <https://slidepeak.com/blog/decktopus-ai-review> (11.11.2025).

Auf dem Weg zu einem kritisch-engagierten Ansatz und einem KI-integrierenden Curriculum im BA-Studiengang „Wirtschaftskommunikation Deutsch“

Abstract

In this article, I present changes to the content of the Bachelor's degree programme *German Business Communication* that I consider necessary in the light of AI-driven developments both in science and economic fields. I focus on developing critical thinking when students use Generative Artificial Intelligence (GenAI) in their scientific work. This fits in well with the programme's critical and engaged academic approach. The article sets out with an empirical consideration of AI in academic settings and GenAI application within the framework of the degree programme. Based upon this, a teaching/learning unit on GenAI use in a module of the programme is presented, which was designed and implemented in the winter semester 2025/26. The preliminary results reveal that further considerations are required to expand teaching units on AI tools for specific areas of business communication.

Keywords kritisches Denken, Künstliche Intelligenz, Lehre, Wirtschaftskommunikation

1 Einleitung

Der sechssemestrige Bachelorstudiengang „Wirtschaftskommunikation Deutsch“ (WKD) an der Westsächsischen Hochschule Zwickau (WHZ) verbindet die fachsprachliche mit der wirtschaftswissenschaftlichen Ausbildung. WKD richtet sich sowohl im grundständigen als auch in der Doppelabschluss-Option an Personen mit Deutsch als weiterer Sprache¹ (DaW). So können Interessierte mit einem niedrigen Deutschniveau (Deutsch A2- und Englisch B2-Niveau gemäß dem Gemeinsamen europäischen Referenzrahmen – GeR) das Studium ab dem ersten Semester beginnen, Doppelabschluss-Studierende steigen auf einem B1+-Niveau im dritten oder im fünften Fachsemester ein. Zum Ende des Studiums liegt das Deutschniveau der Studierenden idealerweise auf B2+ oder sogar auf C1 (GeR), insbesondere auch in den Spezialisierungen Marketing, Personalmanagement, Logistik oder Unternehmensführung.

Für die Weiterentwicklung der Studieninhalte im WKD-Curriculum spielt – neben der gegenwärtigen, multiplen Krisensituation – die hochdynamische Entwicklung Künstlicher Intelligenz (KI) eine herausragende Rolle. Besonders die generative KI (genKI) verändert das wissenschaftliche Arbeiten und die beruflichen Anforderungen an die zukünftigen Arbeitnehmenden.

Im Beitrag argumentiere ich deshalb dafür, neben grundlegenden Kenntnissen zur Funktionsweise von KI v. a. die persönlichkeitsfördernde Fähigkeit des kritischen Denkens und der Reflexionsfähigkeit in Bezug auf genKI im Rahmen des Studiengangs zu fördern. Dafür gehe ich wie folgt vor: Nachdem ich in Abschnitt 2 zunächst den Studiengang vorstelle, werde ich auf die theoretischen Grundlagen (Abschnitt 3), d. h. Persönlichkeitsentwicklung und kritisches Denken, *Deskilling* und den Ansatz einer kritisch-engagierten Wirtschaftskommunikation eingehen. In Abschnitt 4 stelle ich eine Lehr-/Lerneinheit vor, die in die Grundlagen und Funktionsweisen von KI einführt, und studentisch ausgewählte KI-Tools für den akademischen Kontext umfasst.

.....

1 Unter Deutsch als weitere Sprache fasse ich sowohl Deutsch als Fremd- und Zweitsprache sowie jegliche andere Sprachkombinationsmöglichkeiten, in denen Menschen Deutsch erwerben und/oder lernen.

Zuletzt gebe einen Ausblick auf Desiderata für die weitere Integration von gen-KI in den Studiengang hinsichtlich potenzieller beruflicher Anforderungen.

2 Der Bachelorstudiengang „Wirtschaftskommunikation Deutsch“ zwischen 2022 und 2025 an der Westsächsischen Hochschule Zwickau

Durch die Besetzung einer halben Professur „Deutsch als Fremd- und Fachsprache mit dem Schwerpunkt Wirtschaftskommunikation im mehrsprachigen Kontext“ und der Entfristung einer vollen Stelle (Lehrkraft für besondere Aufgaben) für die sprachpraktische Lehre erhielt der Studiengang WKD im Wintersemester 2023/24 dringend notwendige Ressourcen und Kapazitäten, um u. a. die Lehre abzusichern, neue Hochschulpartnerschaften aufzubauen und den Studiengang weiterzuentwickeln.

2025 erfolgte eine Herauslösung von WKD aus dem Mantel-Studiengang „Languages and Business Administration“, womit aus dem ehemals siebensemestrigen Studienschwerpunkt ein eigenständiger sechssemestriger Studiengang entstand². WKD erweist sich nun als anschlussfähiger an ähnliche Bachelorstudiengänge und Masterstudiengänge im europäischen Kontext. Zugleich ermöglicht diese Änderung Studierenden europäischer Kooperationshochschulen, den Aufenthalt an der WHZ für den Erwerb des Doppelabschlusses mit einem Erasmus+-Stipendium zu finanzieren. Für Doppelabschlussstudierende aus Drittländern ist die Verringerung der Studienzeit ebenso vorteilhaft, da nur wenige Stipendien zur Verfügung stehen und der Auslandsaufenthalt mehrheitlich privat finanziert wird. So wird der Doppelabschluss mehr Stu-

.....
2 Auf der Homepage findet sich eine detaillierte Beschreibung des Studiengangs „Wirtschaftskommunikation Deutsch“: <https://www.whz.de/spr/studium/wirtschaftskommunikation-deutsch-ba/>.

dierenden zugänglich und entsprechende Kooperationsabkommen sind einfacher umzusetzen.

3 Neue Anforderungen an WKD-Studierende vor dem Hintergrund der digitalen Transformation

Hochschulbildung ist durch die Dimensionen akademische Bildung (Fachwissenschaft), Arbeitsmarktvorbereitung und Persönlichkeitsbildung geprägt (Wissenschaftsrat 2015: 39). Diese drei Aspekte sind in allen Studiengängen präsent und sollten idealiter in einem ausgewogenen Verhältnis zueinander stehen. Als Hochschule für Angewandte Wissenschaften sind die Studiengänge der WHZ generell durch eine hohe Anwendungsorientierung und die Nähe zur Wirtschaft charakterisiert. Die im Beitrag fokussierte Dimension der Persönlichkeitsbildung liegt quer zu den anderen beiden Dimensionen und ist mit den Fach- und Methodenkompetenzen eng verwoben (vgl. Leiber 2016: 9).

Aktuell erachtet eine Mehrheit der in der Bitkom-Studie befragten Unternehmen (81 %) KI als die wichtigste Zukunftstechnologie; damit ändert sich auch das von Arbeitnehmenden erwartete *Skill-* (und *Mind-*)*Set* (vgl. Bitkom Research 2025). Dem *Future of Jobs Report* nach wird sich die Relation von menschlicher und maschineller Arbeit bis zum Jahr 2030 zugunsten eines höheren Anteils der Zusammenarbeit zwischen Menschen und Maschine (hybride Arbeit) bei einer gleichzeitigen Verringerung von 50 % auf 37 % der rein menschlich verrichteten Aufgaben verschieben (vgl. World Economic Forum 2025).

Insbesondere in den von unseren Studierenden oft gewählten Studienschwerpunkten Marketing, Personalmanagement und Logistik stellen KI-getriebene Automatisierungsprozesse einen wichtigen Faktor dar, durch den sich Wettbewerbsvorteile für die Unternehmen ergeben können. So sorgt der kompetente Einsatz von KI-Tools im Onlinemarketing für ein besseres Verständnis des Kundenverhaltens, da effizientere Analysen großer, dynamischer Daten-

mengen in Echtzeit realisiert oder Kundenerfahrungen durch automatisierte und personalisierte E-Mail-Kampagnen optimiert werden können. Im Personalwesen stellen z. B. das Bewerber:innenmanagement, die Optimierung von Stellenanzeigen oder der Einsatz von Chatbots für Fragen von Kandidat:innen KI-Anwendungsbeispiele dar. Die smarte Logistik arbeitet umfassend mit intelligenten Technologien und teilautonomen Systemen.

All diese Tätigkeitsbereiche erfordern ein hohes Sprachniveau im gemein- und fachsprachlichen Bereich, eine breite rezeptive und produktive Textsortenkompetenz sowie ein umfängliches Fach- und Weltwissen. Zudem wird immer wieder betont, dass Studierende eine kritische Haltung bzgl. KI-bezogener Arbeitsprozesse und KI-generierter Arbeitsergebnisse ausbilden.

3.1 Beobachtungen zum Einsatz von genKI durch unsere Studierendengruppen beim wissenschaftlichen Arbeiten

Das kritische Denken kann in der Hochschule gut durch das wissenschaftliche Arbeiten gefördert werden (vgl. Kruse 2024: 212–230). Folgender Abschnitt geht auf Beobachtungen zum wissenschaftlichen Schreiben und Lesen unserer Studierenden ein.

In den fachsprachlichen Deutsch- und Englischkursen im Studiengang WKD bemerken wir seit der öffentlichen Verfügbarkeit von ChatGPT im November 2022, dass die Studierenden häufig auf KI-Werkzeuge für die Textrezeption zurückgreifen, z. B. für die Zusammenfassung deutscher Fachtexte oder ihre Übersetzung in die Erstsprache der Lernenden (im Seminar). So werden Aufgaben zwar in Sekundenschnelle gelöst, aber ein Lerneffekt bleibt meist aus.

Auch für die Textproduktion im Deutschen (Hausarbeiten, Präsentationen, Portfolios etc.) greifen die Studierenden schnell auf KI-Tools zurück. Problematisch ist bei einer an genKI ausgelagerten Textproduktion, dass der aktive Prozess der Bedeutungskonstruktion umschifft wird und ein *knowledge transforming*, das durch iterative Überarbeitungsschleifen erst tiefergehendes Wissen entstehen lässt (vgl. Bereiter/Scardamalia 1987), gar nicht erst statt-

finden kann. Damit wird das didaktische Potenzial von Schreiben und Lesen erstens für (fach)kommunikative Ziele verschenkt, wie z. B. das Textsorten-repertoire mit den erforderlichen sprachlichen Mitteln und fachliches Wissen zu entwickeln, und zweitens für epistemische Ziele, z. B. das Schreiben und Lesen zur Unterstützung von Denk- und Ideenfindungsprozessen zu nutzen, und drittens für identitätsbildende Ziele, wie die schreibende, reflexive, produktive Auseinandersetzung mit eigenen Lernerfahrungen (vgl. auch Kruse 2024).

Weitere Bedenken und Herausforderungen in diesem Kontext sind, dass der vermehrte Rückgriff auf genKI zu einer „metacognitive laziness“ (Yan et al. 2025: 435) sowie zur Selbstüberschätzung eigener Kompetenzen (Lee et al. 2025) führen kann. Problematisch ist zudem, dass Studierende sich mit KI-generierten Texten weniger identifizieren und sich schlechter an deren Inhalte erinnern (Kosmyna et al. 2025). Sie nehmen inhaltliche Positionen KI-generierter Texte weniger als ihre eigenen wahr und verteidigen diese Positionen auch seltener in Debatten (vgl. Stadler et al. 2024).

In unseren Lehrveranstaltungen beobachten wir daran anschließend, dass die Studierenden die KI-generierten Inhalte oft weder hinterfragen noch auf Korrektheit hin überprüfen und die Inhalte häufig nicht tiefergehend verstehen. Dies zeigt sich darin, dass sie keine Verantwortung für ihre Arbeitsergebnisse und deren Inhalte übernehmen: So können sie bspw. auf Nachfragen zu den mit KI generierten Texten oder Bildern nicht mit angemessenen Erklärungen antworten oder ihre Analysen und Bewertungen mithilfe von Argumenten verteidigen. Es manifestiert sich auch, dass ihr fachliches Wissen mitunter oberflächlich bleibt, sie es nur schwer auf ungewohnte Kontexte transferieren können und selbst einfache Textsorten dem erwartbaren Deutschniveau nicht entsprechend verfasst werden können.³

3 Es ist klar, dass genKI auch zur Verbesserung und Effizienzsteigerung in vielen Bereichen führen kann (vgl. Deutscher Ethikrat 2023). Wir wissen z. B., dass v. a. schwächere Studierende durch die Verwendung von LLMs bessere Noten beim Verfassen von Hausarbeiten erzielen (vgl. Lee et al. 2025). Für das Sprachenlernen können KI-basierte Sprachsoftware, -plattformen oder Chatbots als kognitives Werkzeug oder Sparringspartner eingesetzt werden (vgl. Döbeli Honegger 2025: 30) und adaptives, orts-, personen- und zeitunabhängiges Lernen besser ermöglichen.

3.2 Wissenschaftliches Schreiben und Lesen als Katalysator für kritisches Denken

Es ist weitgehend bekannt, dass die Komplexität der deutschen Wissenschafts- und Fachsprache(n) Studierende mit DaW vor besondere Herausforderungen stellt. Insbesondere das wissenschaftliche Schreiben gilt als eine „überaus anstrengende mentale Tätigkeit“ (Roll 2021: 273), die viel Zeit in Anspruch nimmt. Wie Buck (2025: 12–15) ausführt, wird das wissenschaftliche Schreiben durch den Einsatz von KI-Tools, wenn es nicht nur zur Abarbeitung der Aufgabe, sondern, wie oft erwünscht, zur vertieften Auseinandersetzung mit fachlichen Problemen eingesetzt wird, nun noch komplexer. Denn KI kann erst dann entlastend sowie für das eigene Denken bereichernd und vertiefend eingesetzt werden, wenn Studierende die Teilaufgaben im Schreibprozess auch ohne Hilfe durch KI erledigen können (vgl. ebd.: 13). Aus der Studienlage sowie unseren Unterrichtsbeobachtungen können wir also vorsichtig ableiten, dass durch ein kognitives Offloading an KI-Werkzeuge von Beginn an sowohl die handwerklichen als auch die noch im Aufbau befindlichen analytischen, kritischen und kreativ-erschaffenden Denk- und Handlungsfähigkeiten sowie die fachlichen Kompetenzen der Studierenden erodieren (*Deskilling*; vgl. Reinmann 2023: 7).

Es war nun schon mehrfach die Rede vom kritischen Denken. Ich möchte im Folgenden knapp mit Kruse (2024: 40–45) eine Definition dessen geben, was ich darunter verstehe. Kruse nennt fünf Aspekte kritischen Denkens, die mir in ihrem Zusammenwirken wichtig im Angesicht von genKI erscheinen: die sorgfältige Prüfung von Behauptungen, Selbstständigkeit, Skepsis, Rationalität sowie ein routinemäßiger sorgfältiger Umgang mit Informationen. Es ist also höchst bedenklich und gefährlich, wenn Menschen ihre intellektuelle Selbstständigkeit preisgeben. So besteht nicht nur die Gefahr, dass die Fähigkeit zum kritischen Denken nicht mehr entwickelt wird, sondern auch, dass soziale Beziehungen leiden und in der Konsequenz Menschen anfälliger für Dogmatismus, Populismus, Fundamentalismus, Ethnozentrismus, Nationalismus und Rassismus werden (vgl. ebd.: 45–50).

Damit ist das hochschulische Ziel der Persönlichkeitsbildung angesprochen, welches die Forderung nach der „Entwicklung einer fachlichen Identi-

tät sowie eines wissenschaftlichen und beruflichen Ethos“ (Wissenschaftsrat 2015: 9) umfasst. Es sollen also Selbst- und Sozialkompetenz bei den Studierenden gefördert und aufgebaut werden, die auf einer kritisch-reflexiven, verantwortungsbewussten und demokratischen, teilhabeorientierten Grundhaltung basieren (vgl. ebd.: 40). Das Ziel ist, dass Studierende als handlungsfähige und autonome Persönlichkeiten auf „sich wandelnde, allenfalls bedingt vorhersehbare Anforderungen vorbereitet sind“ (ebd.: 39).

4 Studieninhalte von „Wirtschaftskommunikation Deutsch“ mit Bezug auf KI-Kompetenzen

Der kritisch-engagierte Ansatz erscheint uns in Anbetracht der gerade diskutierten Ausgangslage angemessen, da er Studierenden ermöglicht, einen verantwortungsbewussten Umgang mit Kommunikation in Bildung und Beruf zu erlangen.

4.1 Kritisch-engagierte Wirtschaftskommunikation und kritisches Denken

Aus engagiert-kritischer Perspektive ist für Janich (2023) die Hinwendung zur Frage nach dem Beitrag von Unternehmenskommunikation zum gesellschaftlichen Wohl zentral. Dem gegenüber steht eine vorrangige Betrachtung der gewinnmaximierenden Optimierung von Unternehmens- und Werbekommunikation aus marketingorientierten, psychologischen Perspektiven heraus (vgl. ebd.). So werden bisher selten unternehmerische Textsorten oder Diskurse aus einem (sozial)kritischen und ethischen Ansatz heraus analysiert

und interpretiert (vgl. Stöckl/Molnar 2018: 266).⁴ Dies sei aber notwendig, da Unternehmen über ihre Kommunikation diskursiv wirkmächtig sind. Über das korporative Agenda-Setting thematisieren sie bestimmte Vorkommnisse und Perspektiven, bestimmen mit, was wir (nicht) über die Hintergründe von Produktionsweisen, Dienstleistungen etc. erfahren. Durch das Framing gestalten Unternehmen die Art und Weise, wie Themen dargestellt werden und beeinflussen dadurch möglicherweise ihre Wahrnehmung (vgl. Huck-Sandhu 2014: 654). Unternehmenskommunikation umfasst also Kommunikate, die als „Praktiken zu behandeln [sind], die systematisch die Gegenstände bilden, von denen sie sprechen“ (Foucault 1988: 74). Folglich wirken ihre Themen und die Darstellung dieser in die Gesellschaft hinein (vgl. Janich 2023: 451). Dies ist insofern wichtig, da in unternehmerischer (Werbe-)Kommunikation z. B. oft problematische *Washing*-Methoden zum Einsatz kommen und Menschen über das soziale, diversitäts- oder umweltbezogene Engagement des Unternehmens hinweggetäuscht werden, um zur Verbesserung der Unternehmensreputation und mit dem Ziel der Profitsteigerung beizutragen.

Es ist erwartbar, dass mit genKI erstellte Kommunikate die in den Trainingsdaten zahlreich vorhandenen misogynen und diskriminierungsrelevanten Inhalte noch verstärkt werden, wenn diese nicht hinterfragt und eine Reproduktion dieser Inhalte verhindert wird. Darum ist es sinnvoll, Janichs (2023) forschungsbezogenen Ansatz didaktisch produktiv zu wenden und davon ausgehend Studierende als zukünftige Arbeitnehmende und Führungskräfte auf diese Unzulänglichkeiten von genKI hin zu sensibilisieren und ihr kritisches Denkvermögen zu fördern. Eine Herausforderung ist, dass die Entwicklung des und das kritische(n) Denken(s) Zeit braucht. Damit liegt es quer zu einer gegenwärtig schnelllebigen Handlungspraxis in Hochschulen und Unternehmen. Dies spricht umso mehr dafür, die Schulung des kritischen Denkens stärker als zuvor in bereits bestehende Module zu integrieren.

.....

4 Es gibt einige Beispiele zu (unternehmerischen) Nachhaltigkeitsdiskursen (Reisigl 2021; Rocco 2017). Andere Arbeiten sind jene von Stöckl/Molnar (2018), die das Genre des *Eco-Advertising* untersuchen oder Rusch (2023), die sich dem Phänomen des Rainbow-/Pinkwashings widmet.

4.2 Der Pilotversuch zu kritisch-engagierter Wirtschaftskommunikation und KI

Für den Pilotversuch zu kritisch-engagierter Wirtschaftskommunikation unter Einbeziehung von KI eignet sich das Modul SPR00920 „Kommunikation und Organisation I im internationalen Kontext“⁵, da hier persönlichkeitsbildende Kompetenzen, wie die kritische Selbstreflexion, die Fähigkeit zur Verantwortungsübernahme und zur angemessenen Kommunikation, eine wichtige Rolle spielen. Weitere Grobziele des Moduls sind, wissenschaftliche Arbeitstechniken zu vermitteln sowie interkulturelle Reflexions- und Handlungskompetenz für Hochschule und Alltag auszubauen. Das Modul wird mit einem Portfolio als Prüfungsleistung abgeschlossen. Dieses Semester sollte das Portfolio vier Einträge aus dem seminarbegleitenden Lerntagebuch der Studierenden sowie die Vorstellung und Reflexion eines von ihnen ausgewählten KI-Tools zum wissenschaftlichen Arbeiten umfassen. Das Verfassen des Lerntagebuchs und des Portfolios sollte nicht mit KI-Tools erfolgen, um Schreiben als Mittel des strukturierenden Denkens zu nutzen und die Inhalte des Seminars, insbesondere den KI-bezogenen Teil, noch einmal in Ruhe zu reflektieren.

Bevor sie in ihrem Portfolio das KI-Werkzeug vorstellen und kritisch reflektieren, sollten die Studierenden das Tool in einer mündlichen Präsentation im Kurs vorstellen und gemeinsam mit ihren Mitstudierenden über dessen Einsatz im hochschulischen Alltag diskutieren

Für die konkrete Umsetzung einer Unterrichtsreihe, die über fünf Einheiten angelegt ist, habe ich mich am AKTIV(ER)-Framework⁶ zur KI-Integration in Bildungssettings von Flick (2025) und an den Empfehlungen von

.....

5 Der Studienablaufplan von „Wirtschaftskommunikation Deutsch“ sowie eine detaillierte Beschreibung des hier betrachteten Moduls können unter Eingabe der Studiengangsnummer 751 eingesehen werden: <https://modulux.whz.de/app-modulux/frontend/studiengaenge/>.

6 Das Akronym AKTIV(ER) steht für **A**usprobieren und **A**nwenden, **K**ommunizieren und **K**larheit schaffen, **T**ransformieren von Lernprozessen, **I**ntegrativ Prüfen und Bewerten, **V**erändern und Mitgestalten, **E**valuieren und Erkenntnisgewinn sowie **R**ückschlüsse ziehen und **R**edesign (vgl. Flick 2025). Das Framework greift auf ein praxisorientiertes Kompetenzraster aus dem schulischen Kontext zurück (vgl. Alles et al. 2024). Dieses umfasst die Kompetenzbereiche Verstehen, Anwenden, Reflektieren und Mitgestalten mit jeweils drei Progressionsstufen (vgl. Falck 2025).

Gimpel et al. (2025) orientiert. Für die fachliche Vorbereitung zu KI-Systemen habe ich u. a. Furbach et al. (2024) konsultiert.

Analyse der Ausgangslage

Es ging zunächst darum, die Ausgangslage zu analysieren und Studierende über ihre Erfahrungen mit KI-Tools im Plenum berichten zu lassen: *Wo begegnet Ihnen im Alltag KI? Welche Erfahrungen haben Sie mit KI bisher gemacht?* Es stellte sich heraus, dass die Studierenden bisher wenig mit KI-Tools an der Hochschule in Berührung gekommen waren. Einige haben das LLM ChatGPT im privaten Bereich verwendet, z. B. zur Inspiration für Kochrezepte im Rahmen einfacher Bedienprinzipien. Die Unterscheidung zwischen schwacher, starker KI und Superintelligenz war neu, ebenso die Tatsache, dass sie alltäglich von schwacher KI umgeben sind und auf diese zurückgreifen, z. B. in Form von Sprachassistenten oder Empfehlungsalgorithmen auf Social Media und Streamingportalen. Dem Einsatz von KI-Tools in akademischen und beruflichen Kontexten standen einige Studierende eher negativ gegenüber. Sie äußerten sich auch nicht auf meine explizite Nachfrage, ob sie fürs wissenschaftliche Schreiben schon einmal LLMs eingesetzt haben. Generell war festzustellen, dass die Studiengruppe zwar einige der aktuell viel diskutierten LLMs kennen, aber kaum etwas über die Funktionsweise von genKI wissen. Ebenso wenig waren Fachbegriffe aus dem Kontext der genKI wie *Prompts* oder *Data Bias* bekannt. Es zeichnete sich ab, dass zum Thema KI bisher nur geringes Wissen in der Gruppe vorhanden war und grundlegende Konzepte und Funktionen vermittelt werden mussten.

Entwicklungsprozesse anstoßen – Wissen über genKI erlangen und Einsatzweisen reflektieren

In der zweiten Veranstaltung wurde Input zu grundlegenden Fachbegriffen und Methoden im Kontext von KI gegeben. Thematisiert wurden Algorithmen, Datenarten, Such- und Schließverfahren, Problemlösung, Maschinelles Lernen (überwacht, unüberwacht, verstärkend), *Deep Learning*, *Generative Pretrained Transformer* (GPT), *Natural Language Processing* (NLP) und Chatbots. Als Hausaufgabe erhielten die Studierenden den Auftrag, sich ein kurzes Video anzusehen, in dem erklärt wird, wie die KI anhand probabilisti-

scher Textverarbeitung Texte, Bilder, Codes oder gar Musik erstellt. Weiterhin sollten sie Informationen zu verschiedenen Problemen und Herausforderungen recherchieren, die im Zusammenhang mit dem Einsatz von KI-Tools diskutiert werden. Außerdem bat ich sie, sich über Prompt-Techniken zu informieren, die fünf wichtigsten LLMs und die dahinterstehenden Unternehmen zu recherchieren sowie die Nachrichten zum Thema KI auf der Themenseite der Tagesschau zu verfolgen.

In der dritten Sitzung wurden zunächst in Kleingruppenarbeit und später im Plenum die Ergebnisse der Studierenden kritisch diskutiert. Es kamen neben Chancen v.a. Risiken beim KI-Einsatz zur Sprache, so z.B. zum Datenschutz und Urheberrecht. Weiterhin beschäftigte die Studierenden der Suizid eines Teenagers in den USA, der mit Unterstützung eines Chatbots von OpenAI Selbstmord begangen hatte. Auch wurden das Prompting kurz angesprochen und verschiedene Techniken vorgestellt. Im Anschluss an die Sitzung habe ich den Studierenden die Promptsammlung zum wissenschaftlichen Arbeiten von Barbara Geyer verlinkt⁷, ein Quiz zu Prompttechniken erstellt sowie weitere Informationen zu Rechtsfragen bei der Nutzung von KI-Tools im Lernmanagementsystem zur Lektüre hochgeladen.

Anwenden und Kommunizieren – Eigene Umgänge mit KI-Tools erproben und kritisch reflektieren

Sowohl im Seminar als auch im Portfolio sollten die Studierenden ihr KI-Tool einer kritischen Reflexion mit Bezug auf die Inhalte, die im Seminar dazu diskutiert wurden, unterziehen. Tatsächlich haben alle Studierenden bei der mündlichen Präsentation kritische Aspekte benannt und nicht nur jene, die bereits vorher zur Sprache kamen. So beklagten sie z.B. ein Fehlen der slowakischen Sprachausgabe des Wissensmanagementtools Notion. Auch die Reihenfolge des Einsatzes von KI, also dass zunächst ein eigener Text verfasst werden sollte, bevor er mit KI-Tools wie DeepL stilistisch verbessert wird, wurde von einer Studentin moniert. Ebenfalls äußerten die Studierenden Bedenken zum hohen Energie- und Wasserverbrauch beim Einsatz und Training

.....
7 Diese ist hier einzusehen: https://padlet.com/innovative_lehre/prompt-sammlung-f-r-wissenschaftliches-arbeiten-8aim1m1ppxxl8vn2.

von KI und betonten die Gefahr von Halluzinationen der KI-Tools für ihre wissenschaftlichen Arbeiten. Zudem kam zur Sprache, dass die Ausgabe von Texten, Quellen etc. nicht in einem neutralen, unpolitischen Raum geschieht und deren Prägungen immer mitbedacht werden sollte.

Inwiefern das Portfolio zu einer tieferen Reflexion der KI-Werkzeuge bei den Studierenden anregt, kann zum Zeitpunkt der Abfassung des Artikels noch nicht abgeschätzt werden. Die Aussagen in der Seminardiskussion verweisen jedoch auf Spuren kritischen Denkens bei den Studierenden, die freilich ausbaufähig sind.

5 Fazit und Ausblick

Der Artikel hatte zum Ziel, erste Ergebnisse der Weiterentwicklung von Inhalten des Studiengangs „Wirtschaftskommunikation Deutsch“ mit Fokus auf die Integration generativer KI-Tools in akademischen und beruflichen Kontexten zu skizzieren. Zunächst wurde der Ansatz einer kritisch-engagierten Wirtschaftskommunikation erläutert und die Notwendigkeit, die oft vernachlässigte Dimension der Persönlichkeitsbildung besser und noch intensiver im Studiengang zu integrieren. Dies lässt sich im Kontext von genKI am Beispiel des wissenschaftlichen Arbeitens gewinnbringend realisieren. Basierend auf Unterrichtsbeobachtungen konnten Ambivalenzen in Bezug auf den Umgang mit genKI durch die Studierenden sowie Risiken des *Deskilling*s herausgearbeitet werden. Weiterhin wurde ein Pilotversuch mit Studierenden nachgezeichnet, der zeigt, wie kritische Denkansätze als Voraussetzung für eigenständige kritische Reflexion bei den Studierenden entwickelt werden können. So wurden die Grundlagen eines kritisch-reflektierten Einsatzes zunächst für das wissenschaftliche Arbeiten gelegt. Für die Zukunft ist es ein Ziel, die Qualität und den Nutzen von KI-Werkzeugen in Bezug auf unternehmenskommunikative Tätigkeiten zu erproben und zu evaluieren. Auf diese Weise können die Studierenden die Strategien des kritischen Denkens weiter entwickeln, was eine Grundvoraussetzung für die erfolgreiche Bewältigung von fachlichen und sprachlichen Herausforderungen im späteren Berufsfeld ist.

Literatur

- ALLES, S. / FALCK, J. / FLICK, M. / SCHULZ, R. (2025): *KI-Kompetenzen für Lehrende und Lernende: Aus der Praxis für die Praxis – eine adaptierbare Basis* [PDF]. Virtuelles Kompetenzzentrum – Schreiben lehren und lernen mit Künstlicher Intelligenz (VK:KIWA).
- BEREITER, CARL / SCARDAMALIA, MARLENE (1987): *The Psychology of Written Composition*. Hillsdale: Lawrence Erlbaum.
- BITKOM RESEARCH (2025): *KI in der deutschen Wirtschaft. Unternehmen zu Einsatz, Chancen, Risiken & Co.* https://www.bitkom.org/Bitkom-Dataverse/KIUnternehmen?tab_name=Hype+oder+Zukunftstechnologie (12.12.2025).
- BUCK, ISABELLA (2025): *Wissenschaftliches Schreiben mit KI*. Tübingen: UKV Verlag.
- DEUTSCHER ETHIKRAT (2023): *Mensch und Maschine – Herausforderungen durch Künstliche Intelligenz. Stellungnahme*. Berlin: Deutscher Ethikrat.
- DÖBELI HONEGGER, BEAT (2025): „Generative Machine-Learning-Systeme. Die nächste Herausforderung des digitalen Leitmedienwechsels.“ *Handbuch Lernen mit digitalen Medien* (3. Auflage). Hrsg. Gerold Brägger / Hans-Günter Rolff. Weinheim: Beltz.
- FALCK, JOSCHA (2025): „KI-Kompetenzen für Lehrende und Lernende.“ <https://joschafalck.de/ki-kompetenzen/> (12.12.2025).
- FLICK, MANUEL (2025): „Das AKTIV-Framework: KI souverän in Schule und Unterricht integrieren.“ <https://www.manuelflick.de/blog/aktiv-framework#analyseausgangslage> (12.12.2025).
- FOUCAULT, MICHEL (1988): *Archäologie des Wissens*. Frankfurt am Main: Suhrkamp.
- FURBACH, ULRICH / KITZELMANN, EMANUEL / MICHAELI, TILMAN / SCHMID, UTE (HRSG.) (2024): *Künstliche Intelligenz für Lehrkräfte. Eine fachliche Einführung mit didaktischen Hinweisen*. Wiesbaden: Springer Nature.
- GIMPEL, HENNER / HALL, KRISTINA / DECKER, STEFAN / EYMANN, TORSTEN / GUTHEIL, NIKLAS / LÄMMERMANN, LUIS / BRAIG, NIKLAS / MAEDCHE, ALEXANDER / RÖGLINGER, MAXIMILIAN / RUINER, CAROLINE / SCHOCH, MANFRED / SCHOOP, MAREIKE / URBACH, NILS / VANDIRK, STEFFEN (2025): „Using Generative AI in Higher Education: A Guide for Instructors.“ *Journal of Information Systems Education* 36.3: 237–256.
- HUCK-SANDHU, SIMONE (2014): „Corporate Messages entwickeln und steuern: Agenda Setting, Framing, Storytelling.“ *Handbuch Unternehmenskommunikation. Stra-*

- tegie – Management – Wertschöpfung. Hrsg. Ansgar Zerfuß / Manfred Piwinger. Wiesbaden: Springer VS: 651–668.
- JANICH, NINA (2023): „Unternehmenskommunikation und Werbung aus der Perspektive einer engagiert-kritischen Wissenschaft. Ein Essay zu einem spannenden Forschungspotenzial.“ *Werbung und PR im digitalen Zeitalter. Grenzen, Übergänge und neue Formate*. Hrsg. Martin Nielsen / Marianne Grove Ditlevsen / Anne Grethe Pedersen. Wiesbaden: Springer VS. 445–467.
- KOSMYNA, NATALIYA / HAUPTMANN, EUGENE / YUAN, YE TONG / SITU, JESSICA / LIAO, XIAN-HAO / BERESNITZKY, ASHLY V. / BRAUNSTEIN, IRIS / MAES, PATTIE (2025): „Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task.“ <https://arxiv.org/abs/2506.08872> (12.12.2025).
- KRUSE, OTTO (2024): *Kritisches Denken und Argumentieren*. (2. Auflage). München: UVK Verlag.
- LEE, HAEKWAN P. / SARKAR, ANIRBAN / TANKELEVITCH, LEONHARD / DROSOS, IOANNIS / RINTEL, STEVE / BANKS, RICHARD / WILSON, MAX (2025): „The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and Confidence Effects from a Survey of Knowledge Workers.“ *Proceedings of the CHI Conference on Human Factors in Computing Systems*: 1–22.
- LEIBER, THEODOR (2016): „Persönlichkeitsentwicklung als elementares Bildungsziel. Methodische Optionen der Umsetzung und Bewertung im Hochschulbereich.“ *die hochschullehre* 2: 1–21. https://www.hochschullehre.org/wp-content/files/diehochschullehre_2016_leiber.pdf (12.12.2025).
- REINMANN, GABI (2023): *Deskilling durch Künstliche Intelligenz? Potenzielle Kompetenzverluste als Herausforderung für die Hochschuldidaktik*. Diskussionspapier Nr.25. Berlin: Hochschulforum Digitalisierung.
- REISIGL, MARTIN (2021): „Weniger ist mehr! Diskurslinguistische Betrachtungen zur Rhetorik der Automobilität.“ *Deutsche Sprache* 4: 352–367.
- ROCCO, GORANKA (2017): „Unternehmerische und städtische Nachhaltigkeitsberichte. Eine textologisch- diskurslinguistische Betrachtung.“ *Nachhaltige Stadtentwicklung. Infrastrukturen, Akteure, Diskurse*. Hrsg. Jens Ivo Engels / Nina Janich / Jochen Monstadt / Dieter Schott. Frankfurt am Main / New York: Campus. 240–260.
- ROLL, HEIKE (2021): „Sprachliche Teilkompetenzen (2): Schriftlichkeit.“ *Handbuch Deutsch als Fremd- und Zweitsprache*. Hrsg. Claus Altmayer / Katrin Biebighäuser / Stefanie Haberzettl / Antje Heine. Berlin: J. B. Metzler. 268–281.

- RUSCH, MICHAELA (2023): „True Colours or Rainbow-Washing Exposed!? – Company Pride in and through Digital and Social Media Reviewed.“ *ILCEA* 51. <http://journals.openedition.org/ilcea/17655> (12.12.2025).
- STADLER, MATTHIAS / BANNERT, MARIA / SAILER, MICHAEL (2024): „Cognitive Ease at a Cost: LLMs Reduce Mental Effort but Compromise Depth in Student Scientific Inquiry.“ *Computers in Human Behavior* 160: 108186.
- STÖCKL, HARTMUT / MOLNAR, SONJA (2017): „Eco-Advertising. The Linguistics and Semiotics of Green(-Washed) Persuasion.“ *The Routledge Handbook of Ecolinguistics*. Eds. Alwin F. Fill / Hermine Penz. New York: Routledge. 259–274.
- WISSENSCHAFTSRAT (2015): *Empfehlungen zum Verhältnis von Hochschulbildung und Arbeitsmarkt – Zweiter Teil der Empfehlungen zur Qualifizierung von Fachkräften vor dem Hintergrund des demographischen Wandels*. Bielefeld. www.wissenschaftsrat.de/download/archiv/4925-15.html (12.12.2025).
- WORLD ECONOMIC FORUM (2025): *Future of Jobs Report 2025*. https://reports.weforum.org/docs/WEF_Future_of_Jobs_Report_2025.pdf (12.12.2025).
- YAN, LIXIANG / GREIFF, SAMUEL / LODGE, JASON M. LODGE / GAŠEVIĆ, DRAGAN (2025): „Distinguishing Performance Gains from Learning when Using Generative AI.“ *Nature Reviews Psychology* 4: 435–436.

Fachdeutsch vermitteln mit KI

Didaktisierung von Lehrmaterialien und neue didaktische Konzepte durch den Einsatz von Künstlicher Intelligenz

Abstract

This article deals with the use of Artificial Intelligence (AI) in teaching German for specific purposes. It examines the potential of AI for the didactic design of exercises and teaching materials, for providing personalized feedback, and for supporting differentiated instruction. As a case study, a project conducted during the summer semester 2025 in Kyrgyzstan is presented, in which ChatGPT-supported exercises were developed and evaluated. Additionally, the perspectives of teachers regarding the opportunities and risks of using AI in specialized language teaching are discussed.

Keywords ChatGPT, Didaktik, Künstliche Intelligenz (KI), Lehrmaterial

1 Einleitung

Der Unterricht in Deutsch als Fachfremdsprache stellt für Lehrkräfte und Lernende eine besondere Herausforderung dar (vgl. Niehaus 2024; Orobchuk/Dyakiv 2024; von Blanckenburg/Derichsweiler 2025). Neben der Vermittlung allgemeiner fremdsprachlicher Kompetenzen müssen die Lehrkräfte fachsprachliche Strukturen, terminologische Präzision sowie komplexe fachliche Inhalte beherrschen. Dies erfordert von den Lehrkräften ein breites Spektrum an Kenntnissen und Fähigkeiten, darunter fremdsprachliche, fachlinguisti-

sche, didaktisch-methodische sowie fachbezogene Kompetenzen (Beier/Möhn 1988).

In den letzten Jahren hat die rasante Entwicklung von Künstlicher Intelligenz (KI) neue Perspektiven für den Fremdsprachenunterricht eröffnet (vgl. Baum/Günay 2025; Hartmann/Möller 2025). Generative Sprachmodelle wie ChatGPT ermöglichen die automatisierte Erstellung von Übungen, die Anpassung an unterschiedliche Sprachniveaustufen sowie die Entwicklung interaktiver Aufgabenformate (vgl. Karataş et. al. 2024: 19344). KI-Sprachmodelle erleichtern auch die Erstellung fachspezifischer Materialien, die bisher mit erheblichem Zeitaufwand verbunden waren. Diese technologischen Entwicklungen werfen jedoch auch didaktische Fragestellungen auf: Welche Qualität weisen KI-generierte Materialien auf? Wie verändern KI-gestützte Werkzeuge die Rolle der Lehrkraft? Und in welchem Ausmaß können solche Tools tatsächlich zur Verbesserung des Fachsprachenunterrichts beitragen?

Vor diesem Hintergrund verfolgt der Artikel das Ziel, das Potenzial von KI-Sprachmodellen für den Fachdeutschunterricht zu untersuchen und ihre Einsatzmöglichkeiten für die Didaktisierung fachsprachlicher Lehrmaterialien zu analysieren. Im Mittelpunkt stehen folgende Fragen:

- (1) Wie kann KI zur Automatisierung und Variation von Übungen in verschiedenen Fachrichtungen und auf unterschiedlichen Sprachniveaus beitragen?
- (2) Wie kann KI individualisiertes Feedback ermöglichen?
- (3) Welche Möglichkeiten eröffnet der KI-Einsatz für eine effektive Binnendifferenzierung?

Der Artikel basiert auf einem Projekt, das im Sommersemester 2025 in Kirgistan am Kirgisisch-Deutschen Institut für Angewandte Informatik durchgeführt wurde. Im Rahmen dieses Projekts wurden mithilfe von ChatGPT grammatische und fachsprachliche Übungen erstellt, analysiert und in ihrer didaktischen Eignung bewertet. Zudem wurde unter den Lehrkräften, die Fachdeutsch unterrichten, eine Umfrage zu ihren Erfahrungen, Erwartungen und möglichen Vorbehalten gegenüber KI-gestützten Lehrmaterialien durch-

geführt. Die Ergebnisse dienen als Grundlage für die Diskussion neuer didaktischer Konzepte zum Einsatz von KI.

2 Fachdeutsch und Anforderungen an Lehrkräfte

Fachdeutschkurse nehmen in Fremdsprachencurricula eine besondere Stellung ein, da sie eng mit fachlichen Kontexten und kommunikativen Anforderungen bestimmter Disziplinen verknüpft sind (Wey 2022: 23). Fachdeutsch umfasst terminologisch präzisen Wortschatz, textsortenspezifische Strukturen, typische Kommunikationsformen (z. B. Berichte, Anleitungen, Präsentationen) und berufsspezifische Handlungsabläufe (Roelcke 2014: 9). Aus solchen sprachlichen und fachlichen Anforderungen entsteht ein komplexes Gefüge, das sowohl für Lernende als auch für Lehrende herausfordernd ist. Beier/Möhn (1988: 135–142) beschreiben vier zentrale Kompetenzbereiche, die Lehrkräfte im fachbezogenen Fremdsprachenunterricht beherrschen müssen:

1. Fremdsprachliche Kompetenz, um linguistische Strukturen der Zielsprache sicher erklären und modellieren zu können.
2. Fachlinguistische Kompetenz, um sprachliche Besonderheiten verschiedener Fachdomänen zu erkennen, zu analysieren und adressatengerecht zu vermitteln.
3. Pädagogische, didaktische und methodische Kompetenz, um Lernprozesse zu steuern, geeignete Materialien auszuwählen und Aufgaben so zu gestalten, dass sie den fachlichen und sprachlichen Bedürfnissen der Lernenden entsprechen.
4. Fachliche Kompetenz: Lehrkräfte sollten grundlegende Kenntnisse der jeweiligen Fachgebiete besitzen, damit die Vermittlung fachsprachlicher Inhalte authentisch und korrekt erfolgen kann.

Ein weiterer zentraler Aspekt ist die Individualisierung des Lernprozesses. Die Vorkenntnisse, Ziele und Bedürfnisse der Lernenden sind heterogen, da

sie häufig aus verschiedenen Studienrichtungen oder beruflichen Kontexten stammen.

3 Künstliche Intelligenz im Fremdsprachenunterricht

Der Stand der Forschung zeigt, dass der Fremdsprachenunterricht bereits in verschiedenen Bereichen von KI-gestützten Anwendungen profitieren kann (vgl. u. a. Wah 2025; Behforouz/Al Ghaithi 2025; Li/Long 2025). Erste Studien¹ belegen, dass KI insbesondere in der Übungsgenerierung, Fehleranalyse und im automatisierten Feedback Potenziale bietet. Im DaF- und Fachsprachenunterricht wurde z. B. untersucht, wie Sprachmodelle beim Schreiben unterstützen, Lernmaterialien individualisiert aufbereiten oder alternative Formulierungen anbieten können (vgl. Zhang 2024; Muñoz et al. 2025). Auch die Möglichkeit, Aufgaben für unterschiedliche Sprachniveaus automatisch anzupassen, wurde wiederholt hervorgehoben (vgl. Fayyumi et al. 2025). Zudem zeigen Pilotstudien (u. a. Johansson/ Larsson 2024), dass KI bei der Erstellung fachbezogener Texte oder Dialoge hilfreich sein kann, da Sprachmodelle in der Lage sind, terminologische Muster und fachspezifische Sprachhandlungen zu reproduzieren. Dennoch betonen viele Forschende (u. a. Van der Lee/Moorman 2024), dass die Qualität der generierten Inhalte stark von der Ausgestaltung der Eingaben (*Prompts*), vom Modell sowie von der Nachbearbeitung durch Fachlehrkräfte abhängt.

Die Chancen, die generative KI für den Fremdsprachenunterricht bietet, sind vielfältig. Sie ermöglicht eine Individualisierung der Lernprozesse (vgl. Chen et al. 2020), indem Lernende personalisierte Aufgaben und direktes Feedback erhalten. Außerdem kann KI Lehrkräfte bei der Materialentwicklung unterstützen, indem sie abwechslungsreiche Übungen, unterschiedliche Schwierigkeitsgrade oder fachspezifische Textsorten schnell und flexibel erzeugt (vgl. Zawacki-Richter et al. 2019). Ebenso eröffnet KI die Möglichkeit

.....
1 Bspw. Li 2024, Yang et al. 2024 und Neumann et al. 2024.

der Binnendifferenzierung, da Materialien automatisch an unterschiedliche Lernstände angepasst werden können: „AI-driven systems have the potential to provide personalized learning pathways by continuously adapting instructional materials and tasks to match individual learner profiles and progress, enabling real-time differentiation across diverse learner populations“ (Holmes et al. 2019: 82).

Allerdings weist die Integration von KI in den Fremdsprachenunterricht deutliche Grenzen auf. Ein zentrales Problem besteht in der Qualitätskontrolle: Sprachmodelle können ungenaue oder falsche Informationen, sog. *Halluzinationen*, erzeugen, insbesondere in fachsprachlichen Kontexten. „Large language models can produce fluent but factually incorrect or nonsensical outputs (so-called *hallucinations*), which raises important concerns for their reliability and safety in practical applications“ (Bender et al. 2021: 611). Das soll am folgenden Beispiel gezeigt werden:

Beispiel für eine KI-Halluzination im Fachsprachenunterricht

Eine Lehrkraft im Fach Deutsch für Ingenieurwissenschaften lässt ein Sprachmodell folgende Aufgabe lösen:

Prompt der Lehrkraft:

Erkläre auf Deutsch, wie ein Drehmomentwandler funktioniert, und nenne die drei Hauptkomponenten.

Antwort der KI (mit Halluzination):

Ein Drehmomentwandler besteht aus Turbine, Kompressor und Resonanzmodul.

Problem:

Die KI halluziniert ein Bauteil („Resonanzmodul“), das in der Technik nicht existiert, aber plausibel *klingt*. Die korrekten Bestandteile wären Pumpe, Turbine und Leitrad.

Auch Fehler in Grammatik, Terminologie oder Logik sind möglich und erfordern eine fachliche Überprüfung durch Lehrkräfte (vgl. Bender et al. 2021). Das soll am folgenden Beispiel veranschaulicht werden:

Prompt an die KI:

Erkläre auf Deutsch kurz, was eine PCR-Analyse ist.

Antwort der KI (mit Fehlern):

Die PCR-Analyse ist ein Verfahren, bei dem Proteine im Blut erhitzt werden, damit sie sich schneller teilen. Danach kann man mit einem Mikroskop sehen, wie viele Virus-Zellen entstanden sind.

Fehleranalyse:

1. Terminologischer Fehler:

- PCR arbeitet mit DNA oder RNA, nicht mit *Proteinen*.

2. Fachlicher/logischer Fehler:

- Proteine „teilen“ sich nicht.
- Viren „entstehen“ nicht durch Erwärmung.

3. Methodischer Fehler:

- Ergebnis der PCR wird nicht im Mikroskop sichtbar gemacht, sondern durch Fluoreszenz- oder Mengenmessung.

Zudem können Fragen des Daten- und Urheberrechtsschutzes sowie ethische Aspekte wie Transparenz, Bias oder Abhängigkeit von Automatisierung problematisch sein.

4 Methodik und Projektbeschreibung (Sommersemester 2025, Kirgisistan)

Im Sommersemester 2025 wurde in Kirgistan ein Projekt zur Vermittlung von Fachdeutsch im Bereich Informatik durchgeführt. Zielgruppe waren 25 Studierende mit Sprachniveaus von A1 bis B2 einer privaten kirgisischen Universität. Der Unterricht konzentrierte sich auf Grammatik- und Übungsaufgaben, wobei das selbstentwickelte Lehrwerk „Deutsch für Studierende der Informatik“ eingesetzt wurde. Das Projekt hatte zum Ziel, die Studierenden mit der Fachsprache Informatik vertraut zu machen, ihre Selbstlernkompetenz zu fördern und die Teilnahme an deutschsprachigen Informatikkursen, z. B. im Rahmen von Auslandsaufenthalten oder Hochschulkooperationen, zu ermöglichen.

Im Rahmen des Projekts wurden die Lehrmaterialien mithilfe von ChatGPT erstellt. Dazu wurden gezielte Prompts formuliert, die Sprachniveau, Fachgebiet (Informatik) und Übungstyp (z.B. Lückentexte, Transformationsaufgaben, Redemittel) vorgaben. Die KI-generierten Übungen wurden anschließend qualitativ analysiert und auf sprachliche Korrektheit, fachliche Präzision, didaktische Struktur und Niveaustufen geprüft. Ziel war es, den Studierenden sowohl die Fachsprache Informatik praxisnah zu vermitteln als auch ihre Selbstlernkompetenz zu stärken, z. B. durch Blended-Learning-Aufgaben und eigenständige Online-Recherche. Ein Beispiel für eine von ChatGPT generierte Übung:

Übungen zum Schreiben von W-Fragen

Stellen Sie W-Fragen.

Beispiel: *Der Computer verarbeitet die Werte.*

→ Was **verarbeitet der Computer**?

Das Rechenwerk bearbeitet den Befehl.

→ Was **verarbeitet das Rechenwerk**?

1. Wenn man einen Computer aufmacht, sieht man verschiedene Bauteile.
2. Zu den Computer-Bauteilen gehören Hauptprozessor, Speicher, Festplatte und Grafikkarte.
3. Auf der Hauptplatine sind alle wichtigen Teile des Computers angesteckt.
4. Man kann auf dem Computer Daten und wichtige Informationen speichern.
5. Zum Speichern braucht man den Hauptprozessor und den Speicher.
6. Der Arbeitsspeicher ist das Kurzzeitgedächtnis des Computers.
7. Die Festplatte ist das Langzeitgedächtnis des Rechners.
8. Auf dem Speicher sieht man Speicherchips und Steckverbindungen.

5 Ergebnisse zum Einsatz von KI für Übungen zur Binnendifferenzierung

Die Analyse der KI-generierten Übungen zeigte, dass ChatGPT zuverlässig Aufgaben für unterschiedliche Sprachniveaus (A1–B2) erstellen konnte, die didaktisch sinnvoll strukturiert waren, eine klare Aufgabenlogik zeigten und den angestrebten Niveaustufen entsprachen. Lückentexte, Transformationsaufgaben und praxisnahe Redemittel für Dialogsimulationen waren sprachlich kohärent und größtenteils fachlich korrekt, wobei bei spezialisierten Informatikthemen gelegentlich Terminologiefehler auftraten. Die Progression von geschlossenen zu offenen Übungen erwies sich als didaktisch sinnvoll und unterstützte die Entwicklung selbstständiger Fachsprachenkompetenz. Die fachsprachliche Genauigkeit variierte jedoch je nach Fachgebiet. Auffällig war, dass die KI schnell große Übungsvolumina erstellen konnte, eine fachliche Nachkontrolle dennoch immer notwendig bleibt.

Durch die Variabilität der KI-Ausgaben konnten Lehrkräfte Übungen für Lernende mit unterschiedlichem Kenntnisstand bereitstellen. Dies zeigte das Potenzial von KI, heterogene Lerngruppen effizient zu unterstützen, wie das folgende Beispiel verdeutlicht:

Aufgabe für schwächere Studierende (B1–)

Aufgabenformat: Stark gelenkte Reproduktions- und Transformationsaufgabe

Lesen Sie den folgenden Pseudocode. Ordnen Sie die passenden deutschen Fachbegriffe zu und ergänzen Sie die Lücken im Text.

Begriffe: Variable – Schleife – Bedingung – Ergebnis

Zuerst wird eine _____ definiert. Danach prüft das Programm eine _____.

Wenn sie wahr ist, läuft die _____. Am Ende gibt das Programm das _____ aus.

Sprachliche Unterstützung:

Wortliste mit Übersetzung (Deutsch–Russisch/Kirgisisch optional)

Beispielsatz: *Eine Variable speichert einen Wert.*

Differenzierungsmerkmal: reduzierte Textproduktion, klare sprachliche Stützen, Fokus auf Reproduktion und kontrollierte Anwendung.

Aufgabe für stärkere Studierende (B1+)

Aufgabenformat: produktionsorientierte, komplexere Aufgabe
Beschreiben Sie den Algorithmus in eigenen Worten (ca. 120–150 Wörter).

Erklären Sie:

- welche Variablen verwendet werden,
- wie die Schleife funktioniert und
- warum der Algorithmus effizient ist.

Verwenden Sie dabei mindestens drei der folgenden Konnektoren: *zuerst – anschließend – deshalb – während – am Ende*.

Zusatzanforderung (optional):

- Kurze mündliche Erklärung (2 Minuten) für eine Präsentation.

Differenzierungsmerkmal: höhere kognitive Anforderungen, freie Textproduktion, stärkere fachsprachliche und kommunikative Kompetenz.

Die Aufgabe ist binnendifferenziert gestaltet: Sprachlich schwächere Studierende bearbeiten eine stark gelenkte Reproduktions- und Transformationsaufgabe, bei der sie anhand eines vorgegebenen Pseudocodes zentrale Fachbegriffe wie *Variable*, *Bedingung*, *Schleife* und *Ergebnis* zuordnen und in einen Lückentext einsetzen. Durch sprachliche Unterstützungsangebote (wie eine Wortliste mit Übersetzung sowie einen Beispielsatz) wird die Textproduktion reduziert und der Fokus auf das Wiedererkennen und kontrollierte Anwenden grundlegender fachsprachlicher Strukturen gelegt. Stärkere Studierende hingegen bearbeiten eine produktionsorientierte Aufgabe mit höherem kognitivem Anspruch, indem sie denselben Algorithmus in eigenen Worten beschreiben, dessen Funktionsweise erklären und seine Effizienz begründen. Die zusätzliche Vorgabe zur Verwendung von Konnektoren sowie eine optionale mündliche Präsentation fördern dabei gezielt die fachsprachliche, kommunikative und argumentative Kompetenz.

6 Ergebnisse der Lehrkräftebefragung

Der an die Fachdeutschlehrkräfte gerichtete Fragebogen wurde auf Grundlage zentraler inhaltlicher Dimensionen des KI-Einsatzes im Unterricht entwickelt. Ausgangspunkt war die Differenzierung zwischen wahrgenommenen Chancen und Risiken, da diese in der Forschung als zentrale Einflussfaktoren für die Akzeptanz neuer Technologien gelten. Mehrere geschlossene Items erfassen mit Likert-Skala die Einschätzung der Lehrkräfte zu potenziellen Mehrwerten von KI, insbesondere in Bezug auf Materialentwicklung, Automatisierung von Routineaufgaben und Möglichkeiten zur Binnendifferenzierung. Ergänzend dazu thematisieren weitere Items Risiken, wie Bedenken hinsichtlich der fachlichen Genauigkeit KI-generierter Inhalte, datenschutz- und urheberrechtliche Fragestellungen sowie die Sorge vor einer übermäßigen Abhängigkeit der Lernenden von KI. Spezifische Items erfassen die Haltung der Befragten gegenüber einer kontrollierten und unterstützenden KI-Nutzung, wobei die didaktische Expertise der Lehrkraft als regulierender und gestaltender Faktor berücksichtigt wird.

Die offenen Fragen dienten der Erhebung individueller Perspektiven, indem sie die Befragten nach eigenen Erfahrungen, konkreten Einsatzszenarien sowie notwendigen Rahmenbedingungen für einen verantwortungsvollen KI-Einsatz befragten. Die Kombination aus standardisierten Likert-Items und offenen Fragen zielte darauf ab, sowohl generalisierbare Einschätzungen als auch differenzierte, kontextbezogene Sichtweisen der Lehrkräfte zu erfassen.

Insgesamt schätzen Lehrkräfte den Einsatz von KI überwiegend als chancenreich zur Materialentwicklung, Automatisierung und Differenzierung ein (vgl. Abbildung 1). Gleichzeitig wurden Bedenken bezüglich fachlicher Genauigkeit, Daten- und Urheberrecht sowie der Gefahr geäußert, dass Lernende sich zu stark auf KI verlassen. Die Mehrheit der Befragten sprach sich für eine kontrollierte, unterstützende Nutzung aus, ergänzt durch die didaktische Expertise der Lehrkraft.

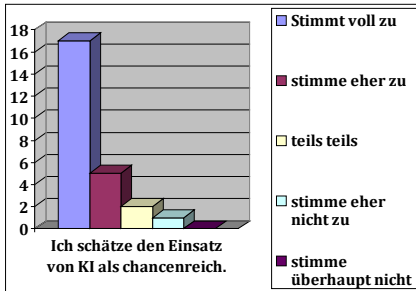


Abb. 1: Chancen von KI (Lehrkräfte)

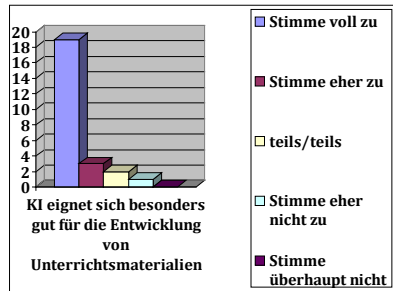


Abb. 2: KI eignet sich besonders gut für die Entwicklung von Unterrichtsmaterialien

Abbildung 2 zeigt auch die Zustimmung, dass sich KI besonders gut für die Entwicklung von Unterrichtsmaterialien eignet. Ablehnende Meinungen („stimme eher nicht zu“ oder „stimme überhaupt nicht“) sind nur in sehr geringem Umfang vertreten.

Abbildung 3 zeigt, dass Lehrkräfte glauben, mit KI sinnvolle Möglichkeiten zur Differenzierung von Lernaufgaben erstellen zu können. Interessant ist, dass die Mehrheit dies allerdings noch nicht ausprobiert hat. Das Ergebnis macht sichtbar, dass viele Befragte das Potenzial sehen, die Nutzung dieser Möglichkeiten jedoch bislang noch nicht erfolgt ist. Abbildung 4 stellt dar, dass die Lehrkräfte Bedenken haben, dass sich Lernende zu stark auf KI verlassen und eigene Lernprozesse vernachlässigen.

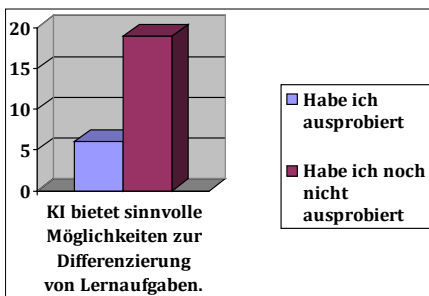


Abb. 3: Differenzierung von Lernaufgaben

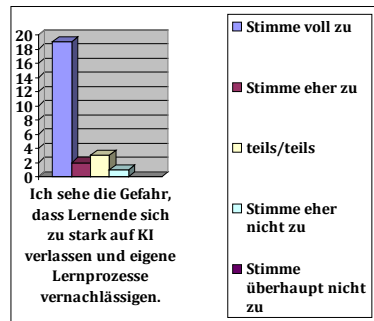


Abb. 4: Datenschutz und Urheberrecht

Bei den offenen Fragen wurden folgende zwei Fragen erörtert:

(1) *Welche Chancen sehen Sie persönlich im Einsatz von KI für Ihren Unterricht? Bitte nennen Sie konkrete Beispiele.*

Die Antworten der Lehrkräfte können wie folgt zusammengefasst werden:

- *Fokus auf Differenzierung und Individualisierung:*
 - „KI ermöglicht es mir, Aufgaben auf unterschiedlichen Niveaus zu erstellen.“
 - „Ich erstelle mit ChatGPT z. B. vereinfachte Texte für lernschwächere Studierende oder Vertiefungsaufgaben für leistungsstarke Lernende.“
- *Unterrichtsvorbereitung und Zeitersparnis:*
 - „Ich nutze KI zur schnellen Erstellung von Aufgaben.“
 - „Ich benutze ChatGPT für die Erstellung von Arbeitsblättern, Quizfragen oder Tafelbildern.“
 - „Mithilfe dieser erstellten Aufgaben habe ich mehr Zeit für die pädagogische Arbeit mit der Lerngruppe.“
- *Unterstützung beim Schreiben und Formulieren:*
 - „KI hilft mir beim Formulieren von Aufgaben oder Texten.“
 - „Ich nutze ChatGPT für die Formulierung von Aufsätzen, Zusammenfassungen oder beim Überarbeiten von Texten.“
- *Förderung selbstständigen Lernens:*
 - „Lernende können mit KI Rückfragen stellen, Erklärungen in eigenen Worten erhalten oder sich Inhalte mehrfach und auf unterschiedliche Weise erklären lassen.“
- *Feedback und Reflexion:*
 - „KI kann erstes Feedback zu Lösungen oder Texten geben (vor allem bei der Erledigung von Hausaufgaben).“
 - „Bevor ich Feedback gebe, kann der Lernende den Lernprozess reflektieren, bevor ich als Lehrkraft gezielt vertieftes Feedback gebe.“

Die Lehrkräfte sehen im Einsatz von KI vor allem Chancen in der Differenzierung und Individualisierung des Unterrichts, etwa durch die Erstellung niveau-gerechter Aufgaben sowie durch Unterstützung beim Schreiben, Formulieren und selbstständigen Lernen der Studierenden. Zudem wird KI als große Ent-

lastung bei der Unterrichtsvorbereitung wahrgenommen, da sie Zeit spart und erstes Feedback ermöglicht.

(2) *Welche Bedingungen müssten aus Ihrer Sicht erfüllt sein, damit KI sinnvoll und verantwortungsvoll im Unterricht eingesetzt werden kann?*

- *Transparenter und reflektierter Einsatz von KI:*
 - „Die Studierenden sollten verstehen, wie KI funktioniert, welche Grenzen sie hat und wie Ergebnisse kritisch zu hinterfragen sind.“
 - „Die Studierenden dürfen nicht alles durch die KI erledigen lassen.“
 - „Die Studierenden sollten die Hausaufgaben nicht mit ChatGPT erledigen lassen.“
- *Technische Ausstattung und verlässlicher Zugang:*
 - „Eine stabile digitale Infrastruktur sowie ein freier Zugang zu geeigneten KI-Tools für alle Lernenden sind Voraussetzung.“
 - „Vor allem in den Regionen ist das Internet schlecht. Das ist ein großer Nachteil für die Studierenden.“
 - „Einige Studierende haben nur ein Mobiltelefon, aber keinen Laptop oder kein Tablet zur Verfügung. Das erschwert ein effektives Lernen und auch eine effektive Nutzung von KI zum Lernen.“
- *Klare Regeln zur Nutzung im Lern- und Leistungsprozess:*
 - „Es braucht verbindliche Absprachen an der Universität darüber, wann und wie KI erlaubt ist.“
 - „Wie sorgen wir dafür, dass die Studierenden selbstständig arbeiten?“
 - „Ich weiß, dass es in Deutschland immer eine ‚Eigenständigkeitserklärung‘ gibt. Bei uns gibt es das nicht.“
- *Verlässlichkeit der Ergebnisse:*
 - „Ich habe schon einige Male bemerkt, dass ChatGPT Fehler macht, zum Beispiel, wenn es um Übersetzungen geht. Das darf nicht sein!“
 - „Ich versuche zu googeln, ob das, was ChatGPT mir gegeben hat, richtig ist.“

Ein sinnvoller und verantwortungsvoller Einsatz von KI im Unterricht setzt voraus, dass Studierende die Funktionsweise und Grenzen von KI verstehen, sie reflektiert nutzen und es klare Regeln zur Verwendung im Lern- und Leistungsprozess gibt.

7 Fazit und Ausblick

Die Analyse der KI-generierten Übungen zeigt, dass generative Sprachmodelle wie ChatGPT grundsätzlich in der Lage sind, fachsprachliche Aufgaben zu erstellen, die sprachlich kohärent und didaktisch sinnvoll sind. Besonders im Hinblick auf Automatisierung und Variation bieten KI-Tools einen erheblichen Mehrwert. Die Ergebnisse der Lehrkräftebefragung bestätigen, dass diese Technologien nach Meinung der Lehrkräfte die Unterrichtsvorbereitung erleichtern und die Individualisierung von Aufgaben unterstützen können. Gleichzeitig wird deutlich, dass die fachliche Präzision und die Qualität zu spezifischen Themenbereichen überprüft werden müssen. Dies entspricht den bisherigen Forschungsergebnissen, die ebenfalls auf die Notwendigkeit einer fachlichen Nachkontrolle hinweisen (vgl. Graf/Fischer 2022; Müller 2023).

Die Integration von KI verändert die Rolle der Lehrkräfte im Fachsprachenunterricht. Sie werden zunehmend zu Moderator*innen, Qualitätsprüfer*innen und Lernbegleiter*innen, die KI-Outputs kritisch prüfen und didaktisch einbetten. Lehrkräfte müssen neue Kompetenzen entwickeln, darunter Umgang mit KI-Tools, Prompting-Techniken und Bewertung der KI-generierten Inhalte.

Die Untersuchung verdeutlicht Chancen und Risiken des KI-Einsatzes: Während Effizienz, Individualisierung und Motivation der Lernenden gestärkt werden, müssen Halluzinationen, fachliche Ungenauigkeiten, Datenschutzfragen und die mögliche Überabhängigkeit der Lernenden sorgfältig berücksichtigt werden. Eine kontrollierte, unterstützende Nutzung von KI ist daher empfehlenswert.

Für die Praxis lassen sich folgende Empfehlungen ableiten:

- KI sollte als ergänzendes Werkzeug in der Materialproduktion und Übungsgestaltung eingesetzt werden.
- Lehrkräfte sollten in der Nutzung von KI geschult werden, insbesondere in Prompting-Techniken und der Qualitätskontrolle.
- Übungen sollten immer einer fachlichen und didaktischen Nachprüfung unterzogen werden.

- Der Einsatz von KI kann als Chance für die Individualisierung, Differenzierung und Entwicklung neuer Aufgabenformate genutzt werden.

Abschließend lässt sich festhalten, dass Künstliche Intelligenz den Fachsprachenunterricht nicht ersetzt, aber signifikant bereichern und modernisieren kann, wenn sie gezielt, reflektiert und unter fachlicher Begleitung eingesetzt wird.

Literatur

- BAUSCH, K. / BAUSCH, M. (2021): *Fachsprachenunterricht: Eine Einführung*. Tübingen: Narr.
- BEHFOROUZ, B. / AL GHAITHI, A. (2025): „AI as a Language Learning Facilitator: Examining Vocabulary and Self-Regulation in EFL Learners.“ *SiSAL Journal* 16.3. DOI: 10.37237/202507.
- BEIER, H. / MÖHN, F. (1988): *Fachsprachenunterricht: Grundlagen und Konzepte*. Tübingen: Narr.
- BENDER, E. M. / GEBRU, T. / McMILLAN-MAJOR, A. / SHMITCHELL, S. (2021): „On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?“ *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT '21)*: 610–623. Association for Computing Machinery. DOI: 10.1145/3442188.3445922.
- VON BLANCKENBURG, M. / DERICHSEILER, S. (2025): „Kompetenzen von Lehrkräften zur Förderung von Streitkultur im Fremdsprachenunterricht.“ *Fremdsprachen Lehren und Lernen* 54: 41–56.
- CHANG, Y. / WANG, X. / WANG, J. / WU, Y. / YANG, L. / ZHU, K. et al. (2024): „Generative AI and Large Language Models – Benefits, Drawbacks, Future and Recommendations.“ *Procedia Computer Science* 246: 5458–5468.
- CHEN, L. / CHEN, P. / LIN, Z. (2020): „Artificial Intelligence in Education: A Review.“ *IEEE Access* 8: 75264–75278. DOI: 10.1109/ACCESS.2020.2988510.
- FAYYOUMI, K. A. R. / YASSIN, B. / ALMARASHDI, H. / AL-KHAWALDEH, M. / AL-FAYYOUMI, S. K. (2025): „The Extent to which Language Teachers Use Linguistic Artificial Intelligence Applications: An Analysis Based on Selected Variables.“ *Front. Educ.* 10:1683845. DOI: 10.3389/educ.2025.1683845.

- FENG, S. / HEFFERNAN, N. T. / KOEDINGER, K. (2022): „Modeling and Leveraging Student Dialogue Behavior with AI in Language Learning.“ *Educational Technology Research and Development* 70.2: 655–679.
- GRAF, M. / FISCHER, S. (2022): *Künstliche Intelligenz im Fremdsprachenunterricht: Chancen und Grenzen*. Berlin: Springer.
- HARTMANN, D. / MÖLLER, M. (2025): „Future Skills und KI-Kompetenz im Sprachunterricht.“ *Künstliche Intelligenz in DaF/DaZ*. (DaF/DaZ in Forschung und Lehre 7.) Hrsg. L. Baum / G. Günay. Berlin: Frank & Timme. 135–151.
- HOLMES, W. / BIALIK, M. / FADEL, C. (2019): *Artificial Intelligence in Education: Promises and Implications for Teaching and Learning*. Boston, MA: Center for Curriculum Redesign.
- JOHANSSON, S. / LARSSON, L. (2025): „Leveraging ChatGPT for Language Learning: A Study on Foreign Teachers’ Swedish Acquisition.“ *Researching Generative AI in Applied Linguistics*. Eds. C. A. Chapelle / G. H. Beckett / B. E. Gray. Iowa State University Digital Press. 40–54. DOI: 10.31274/isudp.2025.211.03.
- KARATAŞ, F. / ABEDI, F. Y. / OZEK GUNYEL, F. / KARADENİZ, D. (2024): „Incorporating AI in Foreign Language Education: An Investigation into ChatGPT’s Effect on Foreign Language Learners.“ *Education and Information Technologies* 29: 19343–19366. DOI: 10.1007/s10639-024-12574-6.
- LI, C. / LONG, J. (2025): „Comparing AI-Assisted and Traditional Teaching in College English: Pedagogical Benefits and Learning Behaviors.“ *Information* 16.10: 895. DOI: 10.3390/info16100895.
- LI, M. (2024): „Leveraging ChatGPT for Second Language Writing Feedback and Assessment.“ *International Journal of Computer-Assisted Language Learning and Teaching* 14.1. DOI: 10.4018/IJCALLT.360382.
- MÜLLER, A. (2023): *Digitale Tools im Fachsprachenunterricht: Eine empirische Studie*. Frankfurt am Main: Lang.
- MUÑOZ MUÑOZ, B. C. / NASSAJI, H. / BELLO CARRILLO, F. I. (2025): „ChatGPT-generated versus Human Direct Corrective Feedback on L2 Writing.“ *System* 134: 103805. DOI: 10.1016/j.system.2025.103805.
- NEUMANN, K. / KUHN, J. / DRACHSLER, H. (2024): „Generative Künstliche Intelligenz in Unterricht und Unterrichtsforschung – Chancen und Herausforderungen.“ *Unterrichtswissenschaft* 52: 227–237. DOI: 10.1007/s42010-024-00212-6.

- NG KOK WAH, I. (2025): „Artificial Intelligence in Language Learning: A Systematic Review of Personalization and Learner Engagement.“ *Forum for Linguistic Studies* 7.9: 327–341. DOI: 10.30564/fls.v7i9.10336.
- NIEHAUS, K. (2024): „Ein Fall für die Förderschule?!“ – Zum Umgang mit sprach(bildungs)bezogenen Herausforderungen.“ *die hochschullehre* 10.1: 319–333. DOI: 10.3278/HSL2427W.
- OROBCHUK, D. / DYAKIV, K. (2024): „DaF-Lehrkräfte an ukrainischen Schulen in Kriegszeit: Herausforderungen und Transformationen der beruflichen Identität.“ *KONTEXTE: Internationales Journal zur Professionalisierung in Deutsch als Fremdsprache* 2.1: 71–88. DOI: 10.24403/jp.1360051.
- ROELCKE, T. (2014): *Fachsprachen*. 3., überarbeitete Auflage. Berlin: Schmidt.
- VAN DER LEE, J. / MOORMAN, K. (2024): „Prompt-based Generative AI in Education: Analyzing the Impact of Prompt Engineering and Human Revision on Output Quality.“ *Journal of Educational Technology & Society* 27.2: 58–71.
- VLČKOVÁ, I. (2023): „Nutzung der künstlichen Intelligenz im Fremdsprachenunterricht.“ *ACC Journal* 29.3: 124–136. DOI: 10.2478/acc-2023-0020.
- WAH, J. N. K. (2025): Artificial Intelligence in Language Learning: A Systematic Review of Personalization and Learner Engagement. *Forum for Linguistic Studies* 7.9: 327–341. DOI: 10.30564/fls.v7i9.10336.
- WEY, S. (2022): „Zur Bedeutung von Sprache im Fachunterricht.“ *Wie Sprache dem Verstehen hilft*. Wiesbaden: Springer VS. DOI: 10.1007/978-3-658-36038-2_2.
- YANG, H. / GAO, C. / SHEN, H. (2024): „Learner Interaction with, and Response to, AI-programmed Automated Writing Evaluation Feedback in EFL Writing: An Exploratory Study.“ *Education and Information Technologies* 29: 3837–3858. DOI: 10.1007/s10639-023-11991-3.
- ZAWACKI-RICHTER, O. / MARÍN, V. I. / BOND, M. / GOUVERNEUR, F. (2019): „Systematic Review of Research on Artificial Intelligence Applications in Higher Education – Where are the Educators?“ *International Journal of Educational Technology in Higher Education* 16: 39. DOI: 10.1186/s41239-019-0171-0.
- ZHANG, Y. (2024): „Incorporating ChatGPT as an Automated Written Corrective Feedback Tool into L2 Writing Class.“ *Journal of Language Teaching* 4.4: 22–34. DOI: 10.54475/jlt.2024.024.

FACHÜBERSETZEN UND DOLMETSCHEN MIT KI-UNTERSTÜTZUNG

JOSÉ TOMÁS CONDE RUANO

Satisfaction with the Use of Microsoft 365 Copilot in Technical Translation Training

Abstract

This chapter explores the satisfaction of third-year technical translation students with Microsoft 365 Copilot in a Spanish university course. Aimed at teachers and researchers in translation training, it examines how generative Artificial Intelligence (AI) may assist specific translation tasks. The methodology involved questionnaires before and after the experience, and following each task, yielding both quantitative and qualitative insights. Results indicate a generally positive attitude towards Copilot, especially in instrumental and strategic tasks, while linguistic and cultural tasks showed limitations. Students reported increased confidence and interest in AI tools, supporting their inclusion in translator training. Findings suggest that Copilot can enhance processes such as organisation and terminology management, though its effective use requires critical guidance. The study underscores the pedagogical potential of generative AI when embedded within ethical and competence-based frameworks.

Keywords Generative AI, Microsoft 365 Copilot, student satisfaction, technical translation, translator training

1 Introduction

The practice of translation has undergone significant transformation in recent years, largely driven by technological advances. As DePalma/Lommel (2025)

argue, the current landscape, shaped by neural machine translation (NMT) and large language models (LLMs), marks both the culmination of decades of evolution and the start of a new paradigm in language transfer.

Many tools that have reshaped the profession — such as computer-assisted translation — emerged from within the language industry, but generative AI represents a turning point. Unlike earlier innovations, LLMs like ChatGPT, DeepSeek or Claude come from broader technological domains, yet they have rapidly permeated professional and academic spheres. Their ability to generate, summarise, and manipulate text, alongside other functions like programming, and image creation, has redefined the boundaries of language-related work.

Translators are among the professionals most directly impacted by this shift. The dual nature of generative AI as both a powerful aid and a potential disruptor has sparked debate across networks, publications and conferences. This study responds to that context. In 2023, the University of the Basque Country adopted Microsoft 365, giving staff access to its full suite, including the generative AI assistant Copilot. This tool enables users to draft, summarise, analyse and automate tasks via natural language. The present research explores its pedagogical potential in technical translation training, assessing student satisfaction across several translation-related tasks.

2 Literature review

Recent studies have examined the performance of generative AI models in intellectual tasks. Dell'Acqua et al. (2023) introduce the concept of the “jagged technological frontier”, showing how AI excels in tasks such as writing and basic analysis, yet struggles with ones that are more complex. Their study with GPT-4 reveals improved quality and speed within this frontier, but reduced accuracy beyond it. Öner/Bengi (2024) note that technical writing with AI demands precise prompting and iterative refinement, requiring new skills aligned with human-machine collaboration (O'Brien, 2023). Xu et al. (2024) find that students appreciate instant feedback from ChatGPT, though its effectiveness depends on user expertise.

The role of generative AI in translation has also been explored. Merali (2024) reports that increased training compute — meaning the amount of processing power and time invested in training LLMs — boosts productivity and quality, especially for less experienced translators. Shahmerdanova (2025) outlines benefits and risks, including linguistic homogenisation, bias and ethical concerns, and advocates a hybrid model combining machine efficiency with human judgement.

There is a growing number of comparative studies between LLMs and NMT systems. Briva-Iglesias et al. (2024) show that GPT-4 and Vicuna outperform Google Translate (GT) in legal translation for several languages. Sun (2024) finds ChatGPT more accurate than DeepL in literary translation. Jiao et al. (2023) confirm that LLMs perform better in major languages but tend to underperform in specialised domains.

Professional translators' attitudes towards generative AI reflect a mix of caution, curiosity and gradual adaptation. Ameri (2024) captures emotional responses from fear to admiration. Rivas Ginel/Moorkens (2024) report growing use of ChatGPT, especially among younger professionals. Common tasks among professional translators include inspiration and terminology searches, with direct translation less frequent (Xu et al. 2025). Concerns include automation and job loss, though many attribute these to production chain dynamics (Jiménez-Crespo 2024). Torrijos Caruda (2022) and Xu et al. (2025) warn of risks such as bias and ethical challenges, especially in culturally sensitive texts.

In higher education, UNESCO (2023) highlights AI's pedagogical versatility, from idea generation to curriculum design. However, challenges related to regulation and academic integrity demand thoughtful implementation. These issues are especially relevant in translation education, where balancing technological support with skill development is essential.

Generative AI introduces new dynamics into translation training. Škobo/Petričević (2023) emphasise the teacher's role in guiding critical use. Jiménez-Crespo (2025) proposes reframing training through the concept of human-centred AI and competence models like EMT and Pym's, which stress instrumental competence.

Empirical studies support this approach. Spoturno (2024) and Tian (2024) stress ethical and critical training, especially in cultural translation. Xu et al. (2025) and Li/Li (2025) note that tools like ChatGPT complement teacher feedback and foster reflection, though uncritical use may reduce lexical richness. They advocate teaching prompt engineering and digital literacy. Li/Liu (2025) explore ethical implications, calling for transparency and collaboration between academia and industry.

Educators' views on AI and NMT have also been studied. Rico/González Pastor (2022) find a general lack of training among teachers, despite positive attitudes towards student use. Sahari et al. (2023) find that teachers prefer GT, while students favour ChatGPT. Both groups recognise benefits but remain cautious about creativity and the translator's role.

Student perceptions of generative AI are increasingly studied. Almulla (2024) finds that ease of use and output quality drive satisfaction. González Pastor (2021) and Aslan (2023) report positive attitudes towards NMT, alongside concerns about professional impact. Tavares et al. (2023) propose activities such as pre-editing and terminology management to foster adaptive competence.

Zhang et al. (2025) find that tools like ChatGPT and Bard support learning, though questions persist about prompt design and dependency. Zhang/Doherty (2025) recommend introducing ethical content early in the degree. Okatan (2025) highlights benefits such as competence development, but also risks like reduced critical thinking. Amaro/Pires (2024) highlight speed and consistency, while Belhassen/Hamda (2025) confirm trust in ChatGPT, though some students question its reliability in complex tasks.

Together, these studies provide evidence in favour of integrating generative AI into translator training, provided it is framed within a critical, ethical and pedagogical approach. Digital literacy, instrumental competence and critical judgement emerge as key pillars for preparing future translators.

3 Materials and methods

This study was conducted with third-year students enrolled in a technical translation course at the University of the Basque Country. Most had previously taken a course with the same lecturer, who had briefly introduced Copilot. The interest shown by the students prompted the present experiment.

Three questionnaires were used: the First Questionnaire (FQ), the Class Questionnaire (CQ), and the Last Questionnaire (LQ). The CQ, shorter in format, was administered five times: once after each given task. Responses were submitted using anonymous codes.

The FQ (including informed consent) assessed students' initial knowledge of AI and attitudes towards its use. It included three demographic questions and six items on knowledge of AI, chatbots and Copilot; prior use of Copilot; perceived usefulness for various tasks; expected translation quality; and agreement with statements on interest, relevance, concerns and ethics. Most items used a four-point Likert scale. The FQ was completed on 23–24 January 2025.

In the next session, students reviewed reading materials and video tutorials on AI, chatbots and Copilot to ensure basic understanding before starting the tasks.

The LQ mirrored the FQ, excluding consent and demographic items. Questions 4, 7, 8 and 9 were repeated, with two new items on the suitability of AI tools in translation classes and across the degree. Questions 5 and 6 were replaced with items on use in other modules, future intentions and broader applications. The LQ was completed on 15–16 May 2025.

The CQ was administered after each task (Table 1). It measured satisfaction and opinions after completing tasks first without and then with Copilot.

Table 1: Overview of competences, assignment and completion date

Competence	Task description	Date
Instrumental	Students receive a translation brief involving the localisation of WordPress templates using Poedit. They first investigate independently how to carry out the task. After 10 minutes, they consult Copilot for guidance.	28/01/25
Linguistic	Students are instructed to produce fluent, idiomatic English-to-Spanish translations of ten sentences. Half the class translates the first five, the other half the remaining five. Once completed, they switch and translate the other set using Copilot.	6–7/02/25
Cultural	Two local Spanish news articles (A and B) are provided. Students are divided into two groups: one translates three culture-specific elements from article A for a British audience; the other does the same with article B. They then repeat the task with the other article using Copilot.	27–28/02/25
Terminological	Students work with two software manuals. Each group creates a mini glossary of three terms from their assigned manual, identifying suitable equivalents. They then repeat the task with the other manual using Copilot.	18/03/25
Strategic	Students simulate a project management scenario for translating a European patent in teams of 7–8 people. They plan the workflow: setting deadlines, dividing tasks, allocating resources. After 10 minutes, they repeat the planning using Copilot.	15/04/25

Tasks were spaced out to prevent anticipation of AI use. The lecturer emphasised completing the initial phase without technological help. The CQ included items on task identification, prompt used, Copilot's role, and a final item on recommending Copilot.

All questionnaires were created using Microsoft 365 Forms and uploaded to the virtual classroom. Students took 10–15 minutes to complete each.

Quantitative data were exported to a spreadsheet; open-ended responses were manually coded using Atlas.ti Web.

4 Results and discussion

Both the FQ and the LQ received 29 responses. Although the respondents may not have been the same — some students dropped the course early while others joined later — the group remained relatively homogeneous.

4.1 Demographic profile

The average age of respondents was 21.4. Of the 29 students, 21 were in their third year, 4 in their fourth, 2 in their fifth, and 2 were exchange students. In response to a question about prior use of AI tools in other courses, eight students mentioned Copilot. More than half had used ChatGPT, while one student cited Perplexity and another Gemini. One referred more generally to “derivatives” of ChatGPT. Five students answered “No”, and five left the question blank.

Common uses of these tools included: understanding or learning (5 mentions), translation (4), summarising or note taking (3), writing (3), information searching (2), and entertainment (1).

4.2 Experience with specific tasks

Student feedback from the CQs revealed varying levels of satisfaction with Copilot across five tasks. The highest ratings were for the instrumental and strategic tasks, with 100% and 91.6% of respondents respectively recommending Copilot. As for the instrumental task, students cited time-saving, clarity and structure as key benefits. However, some noted limitations such as lack of depth (5 mentions), inaccuracy (3), and absence of visual aids. These findings

echo Dell'Acqua et al. (2023), who describe AI's strength in structured tasks but reduced reliability in more complex scenarios. In the strategic task, most comments highlighted useful and well-structured explanations. Yet, three students found Copilot's suggested steps unrealistic.

The terminological and cultural tasks received evaluations that are more mixed. Despite some low ratings, 64.7% and 75% of students, respectively, recommended Copilot. In the cultural task, 68.7% felt Copilot improved their work, though four students criticised its lack of precision. Eight praised its completeness. In the terminological task, precision was the most cited advantage (8 mentions), while two students questioned the reliability of sources. These results align with Shahmerdanova (2025), who warns of bias and inconsistency in culturally sensitive or specialised content.

The linguistic task received the lowest ratings: only 52% found Copilot helpful, and 60% recommended its use. Students often preferred human-generated translations but still saw value in Copilot as a supplementary tool. Some noted that repeated prompting improved its output. Frequent complaints included excessive literalness (16 mentions) and the need to insist for better results (13). This is consistent with Öner/Bengi (2024), who analysed AI-assisted technical writing.

Two additional open questions explored preferences for using or avoiding Copilot and its limitations. Most students repeated earlier comments or left the first question blank. Regarding limitations, most answered "No", though some mentioned the need for repeated prompting (12), lack of naturalness (6), imprecision (6), missing information (3), need for training (3), and poor prompt design (3).

4.3 Evolution of student attitudes

The LQ revealed a clear shift in attitudes. Average scores increased for knowledge of AI (+0.62), chatbots (+0.52), and especially Copilot (+1.62), likely due to its lower initial familiarity.

Students rated Copilot's usefulness across various translation-related tasks. All areas showed improvement, particularly to solve professional queries

(+0.59), technical support (+0.55), and organisation (+0.45). Even translation itself saw a modest rise (+0.31).

Perceived translation quality also improved as for general translation (+0.34), specialised translation (+0.27), and technical translation (+0.14). The smaller increase in technical translation suggests lingering doubts about Copilot's reliability in this domain.

Students expressed greater interest in AI after the experience. Notably, concern about future professional impact (−0.31) and ethical dilemmas (−0.24) decreased. These shifts suggest that classroom engagement helped demystify the technology, in line with findings by Zhang et al. (2025), who report that exposure fosters motivation.

Finally, students strongly supported the experience and the integration of Copilot into translator training. The two items included only in the LQ ("I'm glad we used it in class" and "It should be part of translator education") received high scores (3.48 and 3.14 out of 4, respectively).

5 Conclusion

This study has explored technical translation students' satisfaction with the use of Microsoft Copilot across five tasks representative of core translation competence. The research design combined pre- and post-experience questionnaires with task-specific questionnaires, enabling a view of students' evolving attitudes and a comparative analysis of key aspects of the translation process.

Findings confirm a generally positive response to Copilot, especially in instrumental and strategic tasks, where its clarity, structure and practical usefulness were highly valued. Linguistic and cultural tasks revealed limitations, yet students recognised its potential as a complementary tool. These results resonate with those of Xu et al. (2025), Jiménez-Crespo (2025) and Zhang et al. (2025), who recommend incorporating generative AI into translator training programmes through critical and ethical frameworks.

In technical translation, Copilot appears to support processes like organisation, documentation and terminology management, though its reliability in specialised tasks remains uncertain. Its effective use depends on human expertise and critical guidance.

This study offers insights for educators, showing that generative AI can foster autonomy, critical thinking and instrumental competence when paired with appropriate training. Limitations of the study include the absence of interviews, as they could have enriched the interpretation of the data, and the need to refine links between tasks and sub-competences. Future research could expand task types or explore other textual genres.

In an increasingly technologised professional landscape — where, as Sun/Muñoz Martín (2025) note, technology has become a structural part of translation work — translator training should focus not only on tools, but also on their integration into workflows. The classroom experience appears to have fostered appropriation of AI, reinforcing the idea that future translators will be both users and designers of their own technological solutions.

References

- ALMULLA, MOHAMMED ABDULLATIF (2024): “Investigating influencing factors of learning satisfaction in AI ChatGPT for research: University students perspective.” *Heliyon* 10: 1–18. DOI: 10.1016/j.heliyon.2024.e32220.
- AMARO, VANESSA / PIRES, MANUEL DUARTE JOÃO (2024): “Found in Translation, Lost in Education: Artificial Intelligence’s Impacts on Translation Tertiary Education in Macao.” *Asian Education and Development Studies* 13.4: 269–281. DOI: 10.1108/AEDS-01-2024-0012.
- AMERI, SAEED (2024): “AI-Induced Emotions: A Translator Vulnerability Perspective.” *Journal of Cognition, Emotion & Education* 2.1: 16–28. DOI: 10.22034/cee.2024.435751.1016.
- ASLAN, ERDİNÇ (2023): “Machine Translation: Perception of Translation and Interpreting Students in Turkey.” *Current Trends in Translation Teaching and Learning E* 10.1: 185–216. DOI: 10.51287/cttl20237.
- BELHASSEN, SALEH / HAMDA, ACHWAK (2025): “Translation Students’ Reliance on and Trust in Artificial Intelligence for Successful Translation Projects: Opportuni-

- ties, Challenges, and Implications.” *Arab World English Journal for Translation & Literary Studies* 9.2: 106–119. DOI: 10.24093/awejtls/vol9no2.7.
- BRIVA-IGLESIAS, VICENT / CAVALHEIRO CAMARGO, JOÃO LUCAS / DOGRU, GOKHAN (2024): “Large Language Models ‘Ad Referendum’: How Good Are They at Machine Translation in the Legal Domain?” *arXiv: 2402.07681*. DOI: 10.48550/arXiv.2402.07681.
- DELL’ACQUA, FABRIZIO / MCFOWLAND III, EDWARD / MOLLICK, ETHAN / LIFSHITZ-ASSAF, HILA / KELLOG, KATHERINE C. / RAJENDRAN, SARAN / KRAYER, LISA / CANDELON, FRANÇOIS / LAKHANI, KARIM R. (2023): “Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of AI on Knowledge Worker Productivity and Quality.” *Harvard Business School Working Paper No. 24-013*: 1–44. DOI: 10.2139/ssrn.4573321.
- DEPALMA, DONALD A. / LOMMEL, ARLE (2025): “Transforming Translation: The Evolution and Impact of AI on Language Transfer and Communication.” Sun/Liu/Moratto (2025): 24–40.
- GONZÁLEZ PASTOR, DIANA (2021): “Introducing Machine Translation in the Translation Classroom: A Survey on Students’ Attitudes and Perceptions.” *Tradumàtica* 19: 1–15.
- JIAO, WENXIANG / WANG, WENXUAN / HUANG, JEN-TSE / WANG, XIN / SHI, SHUMING / TU, ZHAOPENG (2023): “Is ChatGPT a Good Translator? A Preliminary Study.” *arXiv: 2301.08745*. DOI: 10.48550/arXiv.2301.08745.
- JIMÉNEZ-CRESPO, MIGUEL ÁNGEL (2024): “Exploring Professional Translators’ Attitudes Towards Control and Autonomy in the Human-Centred AI Era: Quantitative Results from a Survey Study.” *Tradumàtica* 22: 276–301. DOI: 10.5565/rev/tradumatica.400.
- JIMÉNEZ-CRESPO, MIGUEL ÁNGEL (2025): “‘If Students Translate Like a Robot ...’ or How Research on Human-Centred AI and Intelligence Augmentation Can Help Realign Translation Education.” *The Interpreter and Translator Trainer* 19.3–4: 277–295. DOI: 10.1080/1750399X.2025.2542022.
- LI, MENGLU / LI, DECHAO (2025): “Human Expertise vs AI Efficiency: A Comparative Analysis of Student and ChatGPT Post-Editing.” Sun/Liu/Moratto (2025): 100–120.
- LI, BO / LIU, YUAN (2025): “Ethical Issues in Translation Education in the AI Era.” Sun/Liu/Moratto (2025): 210–229.

- MERALI, ALI (2024): "Scaling Laws for Economic Productivity: Experimental Evidence in LLM-Assisted Translation." *arXiv*: 2409.02391. DOI: 10.48550/arXiv.2409.02391.
- O'BRIEN, SHARON (2023): "Human-Centred Augmented Translation: Against Antagonistic Dualisms." *Perspectives* 32.3: 391–406. DOI: 10.1080/0907676X.2023.2247423.
- OKATAN, SEMİH (2025): "Impact of Artificial Intelligence on Translation Education: Benefits of Artificial Intelligence and Challenges Faced by Students." *Istanbul University Journal of Translation Studies* 22: 234–247. DOI: 10.26650/iujts.2025.1611407.
- ÖNER, İŞİN / BENGI, ZEHRA BEĞÜM (2024): "Essential or Obsolete? The Role of Human Competencies in the Tech-Driven Language Services Industry." *transLogos Translation Studies Journal* 7.1: 78–104. DOI: 10.29228/transLogos.66.
- RICO, CELIA / GONZÁLEZ PASTOR, DIANA (2022): "The Role of Machine Translation in Translation Education: A Thematic Analysis of Translator Educators' Beliefs." *Translation and Interpreting* 14.1: 177–197. DOI: 10.12807/TI.114201.2022.A010.
- RIVAS GINEL, MARÍA ISABEL / MOORKENS, JOSS (2024): "A Year of ChatGPT: Translators' Attitudes and Degree of Adoption." *Tradumàtica* 22: 258–275. DOI: 10.5565/rev/tradumatica.369.
- SAHARI, YOUSEF / AL-KADI, ABDU M. TALIB / ALI, JAMAL KAID MOHAMMED (2023): "A Cross-Sectional Study of ChatGPT in Translation: Magnitude of Use, Attitudes, and Uncertainties." *Journal of Psycholinguistic Research* 52: 2937–2954. DOI: 10.1007/s10936-023-10031-y.
- SHAHMERDANOVA, ROYA (2025): "Artificial Intelligence in Translation: Challenges and Opportunities." *Acta Globalis Humanitatis et Linguarum* 2.1: 62–71. DOI: 10.69760/aghel.02500108.
- ŠKOBO, MILENA Z. / PETRIČEVIĆ, VEDRAN D. (2023): "Navigating the Challenges and Opportunities of Literary Translation in the Age of AI: Striking a Balance Between Human Expertise and Machine Power." *DHS* 2.23: 317–336. DOI: 10.51558/2490-3647.2023.8.2.317.
- SPOTURNO, MARÍA LAURA (2024): "Traducción literaria e inteligencia artificial: consideraciones para la formación universitaria." *Cadernos de Tradução* 44.1: 1–26. DOI: 10.5007/2175-7968.2024.e100602.
- SUN, SANJUN / LIU, KANGLONG / MORATTO, RICCARDO (eds.) (2025): *Translation Studies in the Age of Artificial Intelligence*. London: Routledge.
- SUN, SANJUN / MUÑOZ MARTÍN, RICARDO (2025): "Reframing Translation Expertise for the AI Era." *Sun/Liu/Moratto* (2025): 56–72.

- TAVARES, CÉLIA / TALLONE, LAURA / OLIVEIRA, LUCIANA / RIBEIRO, SANDRA (2023): "The Challenges of Teaching and Assessing Technical Translation in an Era of Neural Machine Translation." *Education Sciences* 13: 1–19. DOI: 10.3390/educsci13060541.
- TIAN, XUJUN (2024): "Personalized Translator Training in the Era of Digital Intelligence: Opportunities, Challenges, and Prospects." *Heliyon* 10: 1–13. DOI: 10.1016/j.heliyon.2024.e39354.
- TORRIJOS CARUDA, CARMEN (2022): "Inteligencia artificial y traducción al español. Proyección, riesgos y responsabilidad." *Puntoycoma* 174: 31–38.
- UNESCO (2023): *ChatGPT and Artificial Intelligence in higher education. Quick start guide*. <https://unesdoc.unesco.org/ark:/48223/pf0000385146> (01.12.2025).
- XU, SIMIN / SU, YANFANG / LIU, KANGLONG (2024): "Integrating AI for Enhanced Feedback in Translation Revision: A Mixed-Methods Investigation of Student Engagement." *arXiv: 2410.08581*. DOI: 10.48550/arXiv.2410.08581.
- XU, SIMIN / SUN, YIFENG / SU, YANFANG / LIU, KANGLONG (2025): "Translation Pedagogy in the AI Era: Design and Application of a ChatGPT-Powered Platform for Training and Feedback." Sun/Liu/Moratto (2025): 100–124.
- ZHANG, WENKANG / LI, ALBERT W. / WU, CHENZE (2025): "University Students' Perceptions of Using Generative AI in Translation Practices." *Instructional Science* 53: 633–655. DOI: 10.1007/s11251-025-09705-y.
- ZHANG, YIA / DOHERTY, STEPHEN (2025): "Investigating Novice Translation Students' AI Literacy in Translation Education." *The Interpreter and Translator Trainer* 19: 234–253. DOI: 10.1080/1750399X.2025.2541478.

Entre el entusiasmo y el escepticismo

El uso de la IA generativa para la mejora de la práctica docente y la calidad en la traducción de la divulgación científica

Abstract

Seeking a balance between the polarized positions currently coexisting in academic and professional contexts, this study explores the pedagogical potential of Generative Artificial Intelligence (GenAI) in translating popular science texts. Acknowledging the enduring relevance of the paradigmatic opposition between “apocalyptic” and “integrated” attitudes (Eco 1964), we propose a measured, critically informed approach grounded in university translation teaching practice. Our hypothesis posits that AI can render translation’s quintessential challenges—such as figurative language and cultural references—more visible and intelligible, while also fostering reflective, responsible student attitudes. Through a German – Spanish case study and utilizing the conversation forum as a means of data collection, we highlight the contrastive capacity and stimulating value of AI, provided it is deployed under instructor guidance within quality-oriented tasks. The results are articulated in twelve emerging theses that advocate for a cautious, supervised integration of AI and aim to stimulate pedagogical debate within translation didactics.

Keywords critical AI integration, Generative Artificial Intelligence (AI), pedagogical innovation, popular science translation, translation didactics

1 Introducción

El reciente desarrollo de modelos de inteligencia artificial (IA) generativa ha transformado de manera profunda y acelerada los procesos de producción y mediación lingüística. Entre los ámbitos más directamente afectados se encuentra la traducción automática, cuyo grado de sofisticación ha alcanzado niveles que hubieran resultado impensables hace apenas una década. En general, la práctica de la mediación interlingüística ha confirmado la apreciable calidad que pueden ofrecer estas herramientas, específicamente en la traducción en los ámbitos de la ciencia y la tecnología; y no sólo en la traducción de textos especializados, sino también en la de textos divulgativos científicos y tecnológicos, sobre todo cuando las mencionadas herramientas se emplean como instrumentos de apoyo más que como proveedoras de soluciones definitivas (Bowker/Buitrago Ciro 2019). En el ámbito universitario, esta constatación ha motivado una creciente exploración de su integración en entornos didácticos que nos atrevemos a afirmar que merece entre tanto la consideración de especialidad autónoma dentro del campo de la investigación en traducción e interpretación y su enseñanza (Bernardini/Gaspari 2025).

Sin embargo, el debate sobre su incorporación a la enseñanza no está ni mucho menos cerrado. Lejos de un consenso, la comunidad académica se muestra profundamente dividida: al entusiasmo de quienes ven en estas herramientas un recurso revolucionario se opone el escepticismo de quienes temen un deterioro de las competencias formativas o de la calidad del proceso traductológico. Esta tensión puede interpretarse a la luz de la célebre oposición entre “apocalípticos” e “integrados” formulada por Umberto Eco, que describía dos actitudes paradigmáticas ante las transformaciones culturales provocadas por los medios de comunicación de masas y las nuevas tecnologías: los apocalípticos ven en la cultura de masas el signo de una decadencia general; los integrados, en cambio, la aceptan con entusiasmo y la consideran una expresión positiva de nuestro tiempo (Eco 1964). Más allá de su formulación original, este marco sigue resultando útil para explicar la polarización actual en torno a la IA generativa.

Desde una perspectiva alejada de cualquier dogmatismo, el presente trabajo adopta una posición deliberativa y pragmática. Partimos de nuestra

experiencia real en el aula —con actividades diseñadas y supervisadas para garantizar la pertinencia académica de los ejercicios propuestos sin descuidar la seguridad pedagógica de los estudiantes— y la sometemos al escrutinio de estudiantes avanzados de Traducción e Interpretación. No se trata, pues, de validar acríticamente el uso de la IA en la práctica docente, sino de valorar, desde una experiencia concreta, en qué medida puede actuar como catalizadora de aprendizaje, reflexión crítica y mejora cualitativa en la formación de traductores.

El estudio que presentamos se centra en el uso guiado de modelos de inteligencia artificial generativa para la enseñanza de la traducción de textos de divulgación científica y tecnológica, con una aplicación específica en la combinación lingüística alemán-español. Se parte de una hipótesis simplicísima: que la IA puede ayudar a visibilizar fenómenos traductológicos complejos que suelen pasar desapercibidos para el estudiante (como la referencia cultural o el lenguaje figurado). A partir de esta premisa, el artículo se estructura del siguiente modo: tras una brevísima revisión de los marcos teórico, conceptual y contextual, se describe la experiencia didáctica, desarrollada en un foro de conversación presencial —no confundir con un foro de chat informal, por favor— en dos fases: primero, análisis comparativo de soluciones traductológicas contenidas en traducciones de textos divulgativos (traducciones generadas por un chatbot y sometidas al escrutinio de los participantes sin poseditar); y, luego, discusión monitorizada con los participantes en el foro con el objetivo de averiguar qué experiencia tenían con la utilización de la IA generativa para traducir, cómo valoraban la calidad de las traducciones generadas por la IA y qué opinaban acerca de los pros y los contras del uso de la IA en la formación de traductores. Por último, se presentan los resultados obtenidos y se formula una serie de tesis emergentes con vistas a futuras líneas de investigación y práctica docente.

2 Marcos teórico, conceptual y contextual

La didáctica de la traducción ha experimentado una evolución significativa en las dos últimas décadas, una evolución marcada por la progresiva incorporación de las realidades digitales y por un énfasis ascendente en la autonomía del estudiante, su reflexividad y el *adiestramiento* en una toma de decisiones fundamentada. Los enfoques contemporáneos han situado en el centro del proceso formativo la capacidad del futuro traductor para identificar problemas, valorar opciones y justificar soluciones traductológicas en función del encargo y la situación comunicativa (Kelly 2005; EMT Board 2022). Paralelamente, la corriente socioconstructivista ha subrayado la importancia del aprendizaje situado y del trabajo colaborativo como entornos privilegiados para el desarrollo de la competencia traductora (Hansen-Schirra/Kiraly 2013; Kiraly 2014).

En este contexto, la introducción de herramientas de IA generativa en el aula de traducción no se interpreta como un hecho disruptivo, sino como la consecuencia lógica de la progresiva digitalización de la práctica profesional y el indispensable reflejo de esta realidad en la enseñanza. Con salvedades de naturaleza casi anecdótica, lo que se discute ya no es si deben emplearse las nuevas tecnologías, sino cómo integrarlas pedagógicamente para que contribuyan a la reflexión crítica y no a la automatización acrítica del proceso.

En los entornos académico y profesional actuales, esta cuestión se vincula estrechamente con la noción de calidad en traducción, que, superado el axioma que la encerraba en los estrechos límites de la corrección formal, se concibe como una noción multidimensional que trasciende sobradamente la mera ausencia de errores lingüísticos. Modelos teóricos ampliamente aceptados (House 2015; Moorkens et al. 2018) entienden la calidad desde una perspectiva funcional y pragmática, en la que resultan determinantes el grado de adecuación del texto meta al propósito comunicativo, al género discursivo y a las expectativas de los destinatarios de la traducción. Los enfoques didácticos modernos insisten, además, en la importancia de desarrollar en el alumnado una cierta *conciencia de calidad* que se materialice en la capacidad para comparar, revisar, evaluar y explicar decisiones traductológicas.

En este sentido, la IA generativa constituye una herramienta cuyo uso resulta particularmente fructífero: sus traducciones, aunque a menudo aparen-

temente intachables, suelen exhibir insuficiencias culturales, inconsistencias estilísticas o incluso, ocasionalmente, errores en la reproducción del sentido del texto origen que los estudiantes pueden aprender a detectar, identificar, criticar fundamentadamente y corregir. Su papel en la enseñanza de la traducción, por tanto, no reside en proporcionar un producto final superior, sino en actuar como un espejo contrastivo que estimule la reflexión sobre la calidad del texto meta y el proceso de traducción.

La utilidad de la IA —y la oportunidad de su uso— se manifiesta muy conspicuamente en la traducción de textos divulgativos científicos y técnicos. La divulgación científica y tecnológica, disciplina híbrida donde las haya, presenta desafíos específicos derivados de su doble orientación: rigor conceptual, por una parte, y accesibilidad discursiva, por otra, dos imperativos que no siempre resultan fáciles de conciliar. Las estrategias de reformulación, la metáfora conceptual, los ejemplos culturales y el equilibrio entre precisión y legibilidad son elementos clave de este tipo textual (Myers 2003; Gutiérrez Rodilla 2005; de Semir 2015). Para los estudiantes, el análisis y traducción de esta clase de textos constituye un *modus operandi* idóneo para ejercitarse e ir adquiriendo competencia traductora a través de la necesidad de tomar decisiones de índole pragmática e implementar estrategias de adaptación cultural más allá de la exigencia de reproducir en la lengua de llegada el contenido especializado de los textos origen.

Es innegable que las herramientas de IA generativa muestran un desempeño notable en la traducción de textos divulgativos; no obstante, tienden a neutralizar la metáfora, simplificar recursos estilísticos o reforzar formulaciones poco castizas, frecuentemente anglicismos innecesarios. Paradójicamente, sucede que estos rasgos textuales, por no decir deficiencias, lejos de restar valor didáctico al uso de la IA, abren posibilidades muy prometedoras para la labor de análisis crítico en el aula.

Ahora bien, como apuntábamos más arriba, el indiscutible potencial formativo de la IA generativa convive con riesgos manifiestos que explican el intenso debate que ha suscitado su incorporación al ámbito educativo y, de manera aún más sensible, a la formación de traductores, intérpretes y mediadores interculturales e interlingüísticos.

Numerosos estudios subrayan su potencial para mejorar la calidad del aprendizaje, facilitar la detección de problemas textuales o liberar tiempo para tareas cognitivamente más exigentes (Bowker/Buitrago Ciro 2019; Bernardini/Gaspari 2025). En traducción, su capacidad para producir borradores rápidos, alternativas estilísticas o reformulaciones múltiples permite diseñar fácilmente secuencias didácticas orientadas a la comparación, la evaluación y la edición. Sin embargo, frente a estas posibilidades, no pocas voces de la ciencia, la enseñanza y la práctica profesional alertan sobre riesgos pedagógicos claros: la dependencia excesiva, la aparente autoridad del *output* generativo, la opacidad de los modelos o la eventual erosión de procesos reflexivos esenciales (Vieira/Mesquita 2025). De ahí que nosotros optemos por una propuesta formativa que insista en la necesidad de una integración guiada, en la que el papel del docente se refuerce en lugar de diluirse.

En este marco, el presente estudio se sitúa deliberadamente en una posición intermedia entre el arrebato y la desconfianza, entre el entusiasmo y el escepticismo: ni apocalípticos ni integrados, proponemos un enfoque realista que explore las ventajas de la IA generativa como catalizador didáctico sin desatender las salvaguardas necesarias para preservar la concepción de la traducción como actividad netamente humana y la autonomía crítica del alumnado.

3 Justificación del corpus, caracterización del género, elección del chatbot y enfoque metodológico

Delimitado el marco teórico, procedemos a describir la aplicación práctica que representa, en rigor, el meollo de este trabajo.

Para la experiencia didáctica seleccionamos cinco entregas de la columna *Stimmt's*, de la sección de ciencia —*Wissen*— del prestigioso semanario alemán *Die Zeit*, textos redactadas por el afamado periodista y divulgador científico Christoph Drösser. Nuestro objetivo era examinar cómo la IA generativa reproduce, interpreta o distorsiona contenido lingüístico y rasgos genéricos

y estilísticos —y también analizar en qué medida la confrontación con sus resultados puede favorecer la toma de conciencia crítica por parte de los estudiantes—, y creemos que este corpus se ha revelado especialmente pertinente para ello.

La columna *Stimmt's*, que puede traducirse al castellano como “¿Es cierto?” o “¿Es cierto que...?”, se publica ininterrumpidamente desde 1997 y, con más de un millar de entregas desde entonces, es la columna periodística más longeva de un solo autor en el panorama mediático alemán (Schnabel/Spiewak 2022). En ella, Drösser contesta a preguntas de los lectores de *Die Zeit* sobre, por ejemplo, mitos populares —como si se pueden eliminar las manchas de vino blanco con vino tinto (Drösser 2005) o si los alcohólicos necesitan dosis más altas de narcóticos para la anestesia (Drösser 2025)—; siempre lo hace en clave científica y las respuestas tienen que estar respaldadas al menos por un estudio científico: en ocasiones, el propio redactor lleva a cabo experimentos para poder ofrecer respuestas fidedignas a sus lectores (Drösser 2017). Los *Stimmt's* son exponentes paradigmáticos de una divulgación periodística de carácter conciso, orientada a resolver cuestiones aparentemente simples —por lo general, sobre curiosidades científicas o fenómenos cotidianos— mediante explicaciones claras, fiables y formuladas con un estilo personal e inconfundible. Su periodicidad —casi todas las semanas—, su sólida estructura y la consistencia en su propósito comunicativo permiten considerarlos, siguiendo la teoría del género propuesta por Hatim y Mason, como un microgénero plenamente reconocible dentro del discurso divulgativo: los géneros —en el sentido de *genre*— son “formas convencionales de textos asociados a ocasiones sociales concretas, por ejemplo, el soneto o la receta de cocina” (Hatim/Mason 1995: 304). Vistos así, los *Stimmt's* cumplen holgadamente con los criterios de convencionalidad, estabilidad funcional y unidad pragmática que ambos autores atribuyen a los géneros textuales.

En efecto, estos textos presentan una estructura nítida y bastante estereotipada, compuesta normalmente por cuatro *movimientos*: (1) formulación directa de la pregunta que motiva el texto; (2) legitimación de la pregunta y contextualización de la misma; (3) exposición sintética del trasfondo científico o factual pertinente y explicación científica, histórica o técnica; (4) cierre que ofrece la respuesta final, a menudo matizada, acompañada de una pirueta

conceptual o estilística o de un guiño humorístico. Esta secuencia recurrente facilita la predicción de la progresión informativa y potencia la accesibilidad del contenido, cumpliendo tanto funciones periodísticas (proveer información contrastada en un espacio reducido) como divulgativas (hacer comprensibles materia y conceptos especializados sin sacrificar el rigor) y claramente pragmáticas (satisfacer las expectativas del lector que busca una respuesta breve, autorizada y bien explicada y, de paso, entretenerse con la lectura).

Obviamente, la elección de este microgénero también responde a criterios didácticos. En la formación de traductores, el conocimiento explícito de las convenciones de género, su propósito comunicativo y su organización interna constituye una herramienta esencial para la toma de decisiones. En textos como los *Stimmt's*, la densidad informativa, la condensación retórica y el logrado equilibrio entre precisión y fluidez plantean retos idóneos para la reflexión en el aula. Además, la fuerte impronta autoral de Christoph Drösser —reconocible en su estilo ágil, su tono levemente irónico y su voluntad de transmitir curiosidad intelectual— introduce un componente suplementario de dificultad: la necesidad de preservar la voz del autor en la lengua meta. Este rasgo permite a los estudiantes superar la mera reproducción del contenido de manera formalmente correcta para adentrarse en la gestión del tono y la sensibilidad hacia el *ethos* del autor o, lo que viene a ser lo mismo, avanzar hacia la adecuación estilística, el aspecto que marca la diferencia entre las traducciones simplemente aceptables y las traducciones de calidad.

Para la traducción de nuestros *Stimmt's* nos decidimos, sin tener que pensárnoslo mucho, por el chatbot de OpenAI ChatGPT. En nuestra experiencia, este modelo resulta particularmente adecuado para actividades didácticas e investigaciones en estudios de traducción porque combina una alta sensibilidad contextual con la capacidad de generar lenguaje natural flexible y matizado. Su diseño permite reconocer y procesar encargos de traducción de forma similar a como lo haría un traductor humano, atendiendo al propósito comunicativo, al registro y a las convenciones discursivas del texto meta. Además, ofrece transparencia en la toma de decisiones —al poder explicitar alternativas, justificar elecciones y comentar transformaciones— y una notable adaptabilidad a instrucciones *finas*, lo que facilita la generación de variantes controladas y reproducibles. Estas características lo convierten en una herramienta útil tanto

para el análisis metatraductológico como para la creación de materiales experimentales en contextos académicos. Para la traducción de las columnas con que hemos trabajado, optamos por la versión gratuita de ChatGPT porque es la más utilizada por los estudiantes de Traducción e Interpretación en nuestra universidad y pretendíamos que nuestro *experimento* fuera reproducible con recursos al alcance de todos.

Para el análisis crítico de las traducciones generadas por el chatbot elegido, optamos por un foro de conversación presencial. Esta modalidad permite acceder al razonamiento metacognitivo de los participantes y observar cómo articulan y reformulan sus argumentos en la discusión grupal para procesar, interpretar y evaluar las decisiones traductológicas de terceras personas, en este caso, las decisiones del chatbot (Garrison/Anderson 2010; Muñoz Martín/Olalla-Soler 2022). Asimismo, el formato dialogado favorece la construcción conjunta de conocimiento y el surgimiento de perspectivas críticas que suelen quedar atenuadas en métodos más individualizados, como entrevistas o cuestionarios (Barbour 2013). Esto resulta imprescindible cuando lo que se pretende es recoger opiniones sobre asuntos que no suelen tratarse de manera expresa o sistemática en el aula. Por último, la trazabilidad del intercambio oral, registrada mediante notas de campo, grabación de audio y ulterior transcripción de lo grabado, proporciona un corpus válido para analizar cómo los estudiantes aplican criterios traductológicos y justifican sus elecciones, en línea con los enfoques socioconstructivistas de la pedagogía de la traducción (Király 2014). En conjunto, el foro constituye un entorno idóneo para observar prácticas argumentativas complejas y representativas del razonamiento que se busca promover en la formación avanzada de traductores.

4 El foro

Nuestro foro tuvo lugar el 26 de noviembre de 2025 en la Facultad de Derecho de la Universidad de Granada; la actividad tuvo una duración aproximada de tres horas y fue íntegramente grabada y transcrita. Además de los dos investigadores que firman este artículo, participaron en él cuatro estudiantes

de posgrado, dos mujeres y dos varones —Anna, Elena, Dani y Víctor— de, respectivamente, 27, 24, 28 y 21 años de edad, todos con un dominio completo o muy avanzado del alemán y del español y experiencia en traducción; actualmente, todos están matriculados en el Máster Universitario en Interpretación de Conferencias de la Universidad de Granada.

Para el ejercicio de microcrítica de la traducción previsto, escogimos: *Über den Durst. Stimmt es, dass man sterben kann, wenn man zu viel Wasser trinkt?* (Drösser 2004); *Stirbt man, wenn Luft in die Adern gespritzt wird?* (Drösser 2011); *Verhindert der Genuss von Alkohol Magen- und Darminfektionen?* (Drösser 2014); *Kann man einem Toten die Augen schließen, indem man sanft darüberstreicht?* (Drösser 2016); y *Reicht es, wenn man sich nur mit Wasser wäscht?* (Drösser 2022).

Antes de la celebración del foro, todos los participantes habían analizado y traducido a vista en clase todos los textos con los que íbamos a trabajar. Cinco días antes de reunirnos, se les notificaron por correo electrónico las instrucciones para que pudieran prepararse para la reunión y se les hizo llegar impresa la mayor parte de la documentación con la que trabajaríamos en el foro, consistente en los cinco textos origen que ya conocían con los pasajes concretos que íbamos a comentar sombreados. En el transcurso de la sesión de trabajo, se les proporcionaron, también impresas, las traducciones de los cinco *Stimmt's*.

El trabajo quedó repartido en dos sesiones separadas por un breve recreo:

- En la primera parte, de una duración aproximada de dos horas, comentamos las traducciones de algunos pasajes, traducciones generadas por ChatGPT y no poseditadas. Aunque los conocían bien porque previamente los habían traducido a vista en el aula, pedimos a los participantes en el foro que antes de nuestro encuentro repasaran los cinco *Stimmt's*, que localizaran los problemas de traducción que contenían los pasajes sombreados y que recordaran cómo los habían solucionado o pensarán cómo los resolverían (aunque, al margen de los textos origen, los documentos de trabajo primarios fueron las traducciones generadas automáticamente). Como puede observarse, operamos con mucho volumen de texto (en total, entre textos origen y textos meta, procesamos 2914 palabras repartidas entre diez

documentos): obviamente no podemos comentarlo todo en este artículo, de manera que aquí tendremos que ceñirnos a mencionar los resultados más relevantes. A riesgo de resultar demasiado insistentes, conviene destacar que el carácter innovador de la actividad reside en que las traducciones generadas por el chatbot no se usan de manera vergonzante y casi *clandestina*, como sucede a veces en nuestras aulas, sino de manera explícita y abierta; y no como si la IA representara el papel de instructor o asistente del profesor, sino, por el contrario, como si el chatbot fuera un participante más en los ejercicios al uso en los que se pone en común, se critica y se evalúa el trabajo de los estudiantes.

- La segunda parte, más parecida a los foros habituales (Barbour 2013), consistió en una conversación guiada y moderada con el objetivo de averiguar qué experiencia tenían los participantes con el uso de chatbots de IA generativa, tanto en la práctica de la traducción como en su formación como traductores. Para ello, les preguntamos para qué tareas usaban la IA y cómo (por ejemplo, si sólo pedían al chatbot que tradujera textos o si también lo empleaban para el trabajo de documentación). Además, queríamos saber qué tipo de instrucciones daban a la IA, más allá de ordenarle que tradujera, esto es, si usaban *prompts* más específicos y cuáles. Obviamente, también nos interesaba averiguar cómo valoraban la calidad de las traducciones generadas por ChatGPT, en especial, en comparación con la calidad de las traducciones hechas por seres humanos —lógicamente, con los textos con que trabajamos en la primera parte del foro en el punto de mira pero sin renunciar a comentar otras traducciones—; queríamos saber en qué aspectos valoraban positivamente el rendimiento de ChatGPT y en qué aspectos consideraban que la *máquina* falla o no proporciona resultados óptimos. Por último, conversamos brevemente sobre las ventajas e inconvenientes —o incluso peligros— del uso de ChatGPT en la formación de traductores: queríamos conocer la opinión de los participantes acerca de la influencia de la IA en el proceso de adquisición de los conocimientos y las destrezas necesarias para traducir. Vistos el reconocimiento como traductor

indudablemente eficiente de ChatGPT y la autoridad que le confiere su condición de LLM más popular y probablemente más evolucionado, deseábamos averiguar si pensaban que el desempeño de nuestro chatbot podía tener un efecto desmotivador —o motivador— para los estudiantes en el transcurso de sus estudios. Esta sesión tuvo una duración aproximada de 45 minutos.

4.1 Los resultados del foro: microcrítica de la traducción

A la hora de evaluar el *output* generado por la máquina, en líneas generales, nuestros informantes consideraron las traducciones notables, sobradamente aceptables aunque no excelentes. Éste no es mal resultado si tenemos en cuenta que, en primer lugar, los textos no habían sido poseditados y, en segundo lugar, los estudiantes de posgrado suelen tener bien afilados sus criterios para el ejercicio de la crítica —que es lo que se ensaya todo el tiempo en el aula de traducción— y resultan por consiguiente felizmente incisivos; debemos hacer hincapié empero en que, en lo tocante a la reproducción del sentido, los participantes en el foro prácticamente no pusieron pegas a las traducciones de ChatGPT. De un modo u otro, no hay traducción en la que no pueda valorarse —positiva o negativamente— algún aspecto: todas resultan mejorables en un sentido u otro y ocasionalmente muestran hallazgos que es de justicia mencionar. A continuación, como solemos hacer en nuestra práctica docente, para presentar los resultados del ejercicio de microcrítica en que consistió la primera parte del foro, nos serviremos de tablas —nueve tablas numeradas— con una estructura de dos columnas y tres filas: en la columna de la izquierda quedará consignado el pasaje del texto origen; en la columna de la derecha, el pasaje equivalente del texto meta; debajo de las columnas con textos origen y meta aparecerá una fila con una única celda con los comentarios correspondientes al cotejo de ambos pasajes. Cabe insistir en que estos comentarios son en su mayor parte el reflejo de las consideraciones aportadas por los estudiantes en el foro.

Tabla 1:

Pasaje n.º 1 (Drösser 2004, 2011, 2014, 2016, 2022)	Traducción generada por ChatGPT, 17/11/2025
Stimmt's?	¿Es cierto?
Nada que objetar a la versión de ChatGPT, que de manera consistente fue la misma para el título de la columna en todas las ocasiones. Los participantes en el foro comentaron, no obstante, que hubieran preferido la traducción “¿Es cierto que...”, pese a que ésta impondría restricciones de orden sintáctico y de redacción (y son conscientes de que, en traducción editorial, las traducciones de los títulos suelen ser establecidas por el cliente). En el <i>prompt</i> dado a ChatGPT, se le indicó de manera un tanto obvia que los <i>Stimmt's</i> son textos breves de divulgación científica extraídos del semanario alemán <i>Die Zeit</i> y que formaban parte de un encargo de traducción editorial para la sección de ciencia de un diario generalista español cuyos lectores tienen buenos conocimientos científicos básicos. El encargo concreto fue: “Para el apartado de su sección de ciencia ‘Qué responde la prensa extranjera a las preguntas de sus lectores,’ un diario español desea traducir la siguiente columna del semanario alemán <i>Die Zeit</i>”. También se le indicó que tradujera como si fuera un traductor profesional, teniendo muy en cuenta las particularidades del género en que se inscriben los textos.	

Tabla 2:

Pasaje n.º 2 (Drösser 2011)	Traducción generada por ChatGPT, 17/11/2025
Stirbt man, wenn Luft in die Adern gespritzt wird? ...fragt Jürgen Sander aus Braunschweig.	¿Te mueres si te inyectan aire en las venas? ...pregunta Jürgen Sander, de Braunschweig.
Respecto de este pasaje, los estudiantes elogiaron que el chatbot hubiera optado por “venas” para “Adern” —que, en rigor, se refiere en alemán a <i>todos</i> los vasos sanguíneos—; “venas” es efectivamente una equivalencia menos precisa y más restrictiva pero los participantes entendieron que encajaba mejor por ser una expresión más popular y más breve, rasgos deseables en el título de una columna de carácter divulgativo. También opinaron que, en este pasaje concreto, la distinción entre diferentes tipos de vasos sanguíneos no resultaba demasiado relevante.	

Pasaje n.º 2 (Drösser 2011)	Traducción generada por ChatGPT, 17/11/2025
En lo tocante al nombre de la ciudad alemana, “Braunschweig”, aunque los estudiantes sabían que en español puede nombrarse como “Brunswick”, no tenían clara la vigencia actual del topónimo <i>español</i> y hubo opiniones divergentes; estuvieron de acuerdo, sin embargo, en que resultaría clarificador para el lector español añadir “Alemania” tras el nombre de la ciudad alemana. Por último, aprobaron que el chatbot se dirigiera al lector español en segunda persona del singular en lugar de utilizar el impersonal.	

Tabla 3:

Pasaje n.º 3 (Drösser 2011)	Traducción generada por ChatGPT, 17/11/2025
Wenn im Fernsehen ein Arzt eine Spritze verabreicht, dann geht die Kamera gern nahe ran und zeigt, wie der Mediziner die Spritze „entlüftet“ und eine kleine Fontäne herausspritzt. In Krimis wird das Spritzen von Luft gern als Tötungsmethode verwendet.	Cuando en la televisión un médico pone una inyección, la cámara suele acercarse para mostrar cómo “purga” la jeringuilla y deja salir un pequeño chorro. En las novelas y series policíacas, inyectar aire es un recurso habitual para asesinar a alguien.
Aquí, el sentido del texto origen está reproducido cabalmente y de manera formalmente correcta, así que desde un punto de vista funcional, apenas cabe objetar nada. No obstante, la expresión “Krimis”, que puede entenderse como referencia cultural muy alemana, no está resuelta de manera óptima: el chatbot parece ser consciente de que “Krimis”, más allá de su apariencia transparente y fácil de traducir, requiere un tratamiento especial porque no existe un equivalente en español que englobe todas las manifestaciones —televisivas, cinematográficas, radiofónicas, literarias, etc.— del género negro o policíaco. ChatGPT lo <i>sabe</i> y, de entre las múltiples posibilidades, elige dos manifestaciones concretas, “novelas y series”. El desglose es lícito; pero sucede que el texto origen toma como referencia el universo audiovisual, del que se aparta la manifestación del género “novelas”. Éstos son los aspectos para cuya detección y explicación viene bien tener a mano un profesor humano razonablemente competente.	

Tabla 4:

Pasaje n.º 4 (Drösser 2016)	Traducción generada por ChatGPT, 17/11/2025
Zwei Drittel aller Menschen, das hat ein amerikanischer Hospizarzt einmal akribisch ausgezählt und in einer Fachzeitschrift veröffentlicht, sterben mit geschlossenen Augen.	Un médico de un hospicio estadounidense, que se tomó la molestia de contar casos y publicar los resultados en una revista especializada, determinó que dos tercios de las personas mueren con los ojos cerrados.
En este pasaje, ChatGPT <i>mete la pata</i> a conciencia, con una inusitada acumulación de errores en lo tocante a la reproducción del sentido, aunque esto sólo lo notarán los críticos avezados que <i>se tomen la molestia</i> de confrontar texto origen y texto meta (con la salvedad de que no <i>dice</i> lo mismo que el original, el <i>output</i> es formalmente correcto y no desentona estilísticamente). Pero vamos por partes: en el original, el médico lleva a cabo meticulosamente una investigación basada en el cómputo preciso de quiénes mueren con los ojos cerrados y quiénes lo hacen con los ojos abiertos; y luego publica los resultados de su investigación en una revista especializada. La traducción sugiere que el médico se aplica en la ejecución de su empeño porque lo que hace es trabajoso, tanto el conteo como la publicación, aunque, por otra parte, dado que nadie se había ocupado de ello antes, tampoco se trata de algo de tanta importancia. Aunque Christoph Drösser se caracteriza por un estilo que ocasionalmente podemos calificar de burlón, en este pasaje concreto no hay sombra de ironía. Por último, “Hospiz” y “hospicio” son falsos amigos parciales: el equivalente de “Hospiz” aquí sería “unidad de enfermos terminales” o, mejor, “unidad de cuidados paliativos” (el “hospicio”, en español actual, remite al orfanato u orfelinato).	

Tabla 5:

Pasaje n.º 5 (Drösser 2022)	Traducción generada por ChatGPT, 17/11/2025
Der US-amerikanische Arzt und Journalist James Hamblin, der an der Yale University Public Health lehrt, hat fünf Jahre lang nicht geduscht und weitgehend auf den Gebrauch von Seife und Shampoo verzichtet. Dann hat er ein Buch geschrieben, das im vergangenen September unter dem Titel <i>Natürlich waschen!</i> auf Deutsch erschienen ist.	El médico y periodista estadounidense James Hamblin, profesor de Salud Pública en la Universidad de Yale, pasó cinco años sin ducharse y prescindiendo casi por completo del jabón y el champú. Luego escribió un libro, publicado en alemán el pasado septiembre con el título <i>¡Lavarse de forma natural!</i>

Pasaje n.º 5 (Drösser 2022)	Traducción generada por ChatGPT, 17/11/2025
Todo iba bien en este pasaje hasta que apareció el nombre del libro de James Hamblin: lo primero, si nos ponemos muy puntillosos, es que “<i>¡Lavarse de forma natural!</i>” no es un título alemán (pero, desde su voluntad de comprender el texto, el lector español entenderá lo que quiere decir el traductor y pasará por alto este detalle); lo segundo es que un buen traductor humano que trabaje concienzudamente y sin prisa no actuará resueltamente y sin vacilaciones, como hace el chatbot —esto es, no <i>alucinará</i>—, sino que muy probablemente se documentará para averiguar cuál es el título del volumen de James Hamblin en español, ya que pensará que ése es precisamente el dato pertinente para los lectores del texto meta (en España, el libro se titula <i>Si nuestra piel hablara: La nueva ciencia de la piel</i>).	

Tabla 6:

Pasaje n.º 6 (Drösser 2022)	Traducción generada por ChatGPT, 17/11/2025
Natürlich wäscht auch James Hamblin offensichtlichen Schmutz von seinem Körper ab. Und gerade in Corona-Zeiten achtet er darauf, sich regelmäßig die Hände zu waschen, mit Seife. Ansonsten sieht er äußerst adrett aus und ist nicht zum sozialen Paria geworden – was er aber auch darauf zurückführt, dass er als junger, gebildeter weißer Mann keine großen Akzeptanzprobleme hat, solange er auf ein gepflegtes Aussehen achtet.	Por supuesto, Hamblin sí elimina la suciedad evidente de su cuerpo. Y, especialmente en tiempos de coronavirus, procura lavarse las manos con regularidad usando jabón. Por lo demás, tiene un aspecto muy pulcro y no se ha convertido en un paria social, algo que él mismo atribuye también al hecho de ser un hombre joven, instruido y blanco, que no suele sufrir grandes problemas de aceptación siempre que cuida mínimamente su apariencia.
Aunque este pasaje presenta algunos defectos, no podría afirmarse, en líneas generales, que la traducción sea inaceptable. Es verdad que el chatbot no especifica cómo “elimina la suciedad evidente de su cuerpo” James Hamblin en la primera frase; que no mantiene la distancia que en la segunda frase del texto origen existe entre el núcleo de la oración y el complemento circunstancial de instrumento “mit Seife”, con las consecuencias estilísticas y pragmáticas que esto conlleva; que no reproduce de la manera más apropiada el sentido de “adrett” en la tercera frase;	

Pasaje n.º 6 (Drösser 2022)	Traducción generada por ChatGPT, 17/11/2025
y que introduce en el texto término alguna matización que no se encuentra en el texto origen (como en la adición de “mínimamente” en la última frase); pero todas las fallas enumeradas hasta aquí son disculpables si el corrector o revisor no aplica criterios de absoluta excelencia en su corrección o revisión: el sentido del texto origen está razonablemente reproducido, la traducción es formalmente correcta y el estilo del texto término no se aparta excesivamente del estilo del texto origen. Ahora bien, aun en lenguas con una sintaxis relativamente libre, como el español, el orden de las palabras sí que puede alterar el producto, y en la sucesión de adjetivos con que se califica a Hamblin en el texto origen, “junger, gebildeter weißer Mann”, el más importante para la caracterización de nuestro protagonista y su estatus social y étnico es “weiß”, que, además, es inseparable del sustantivo al que acompaña y el único de la tríada del que no se puede prescindir: en rigor, “jung” y “gebildet” califican en el texto origen a un “weißer Mann”. Por todo ello, “blanco” debería aparecer en primer lugar en la relación de adjetivos de la traducción: no “hombre joven, instruido y blanco”, sino “hombre blanco joven e instruido” (a menos, claro está, que se quiera escamotear el rasgo crítico claramente expresado en el texto alemán).	

Tabla 7:

Pasaje n.º 7 (Drösser 2022)	Traducción generada por ChatGPT, 17/11/2025
Er [Hamblin] rieche „wie ein Mensch“, sagte seine Freundin – ein Geruch, den die meisten von uns heute kaum noch kennen, weil wir peinlich bemüht sind, ihn täglich abzuwaschen und zu überparfümieren.	“Hueles a persona”, le dijo su novia: un olor que la mayoría de nosotros apenas reconocemos ya, porque nos empeñamos en eliminarlo y cubrirlo de perfume cada día.
De manera algo inconsistente, ChatGPT convierte aquí el estilo indirecto del texto origen en estilo directo y altera la persona gramatical (tercera en el texto origen, segunda en la traducción), de manera que lo que en alemán es una aseveración medianamente neutra se transforma en español en una interpelación directa de su novia al protagonista de este <i>Stimmt's</i>. Al margen, en español, en “persona” se incide más en el carácter individual de cada uno de nosotros que en “ser humano”, expresión en la que prima la pertenencia a la especie entendida como colectividad, como precisamente se expresa en alemán con “Mensch”.	

Tabla 8:

Pasaje n.º 8 (Drösser 2014)	Traducción generada por ChatGPT, 17/11/2025
Im Mai 2000 nahmen 120 Menschen an einem Festbankett in der spanischen Stadt Castellón teil. Für mehr als die Hälfte von ihnen endete der Abend wenig feierlich: Sie wurden von Übelkeit und Durchfall geplagt, hatten Magenkrämpfe und Fieber. Wie sich herausstellte, hatten sie sich mit Salmonellen infiziert, die sich in der Mayonnaise eines Kartoffelsalats und in Thunfisch-Sandwichs befanden.	En mayo del año 2000, 120 personas asistieron a un banquete en la ciudad española de Castellón. Para más de la mitad, la velada terminó de forma poco festiva: sufrieron náuseas y diarrea, cólicos abdominales y fiebre. El análisis posterior reveló que se habían infectado con salmonela presente en la mayonesa de una ensaladilla y en unos bocadillos de atún.
Al principio de este pasaje, hay soluciones mejorables que cabe comentar: primero, con lectores meta españoles, no es necesario especificar que Castellón es una “ciudad española”; en segundo lugar, “cólicos abdominales” seguramente no sea la mejor equivalencia para “Magenkrämpfe” —aunque, en la práctica, cuando un paciente dice que tiene “cólicos abdominales”, suele querer decir que tiene “calambres estomacales”, y viceversa—. Ahora bien, donde ChatGPT hace su trabajo particularmente bien, hay que reconocerle su mérito: este pasaje, aparentemente sencillo, tiene su miga. No caben muchas dudas acerca de que el “Kartoffelsalat” y los “Thunfisch-Sandwichs” del texto origen son adaptaciones culturales para que el lector alemán se haga una idea aproximada de lo que comieron los festejantes. Es digno de elogio que el chatbot haya recorrido en sentido inverso el camino andado por el redactor de <i>Die Zeit</i> y haya deshecho la adaptación cultural del <i>Stimmt’s</i> al recuperar los manjares originarios: no “ensalada de patata” ni mucho menos “sándwiches de atún” —que un castellonense de pro apenas consideraría comestibles—, sino, con toda seguridad, “ensaladilla [rusa]” y “bocadillos de atún”. Eso o que, sin plantearse los motivos del redactor del texto origen, el chatbot haya hecho una adaptación cultural del universo culinario alemán al español para que al lector meta no le causen extrañeza la “ensalada de patata” o los “sándwiches de atún” en un banquete en Castellón. De un modo u otro, el resultado es más que notable.	

Tabla 9:

Pasaje n.º 9 (Drösser 2004)	Traducción generada por ChatGPT, 17/11/2025
Über den Durst Stimmt es, dass man sterben kann, wenn man zu viel Wasser trinkt? ... fragt Rita Bartmann aus Frankfurt.	¿Es verdad que uno puede morir por beber demasiada agua? ... pregunta Rita Bartmann, desde Fráncfort.
En este pasaje, ChatGPT sencillamente omite la traducción del título de la columna, “Über den Durst”, que nuestros estudiantes se limitaron a traducir, sin más, como “Sobre la sed” o “Acerca de la sed” ignorando que “Über den Durst” hace referencia a la expresión alemana “einen über den Durst trinken”, que puede traducirse como “beber más de la cuenta” o “beber en exceso”, originariamente en el sentido de beber alcohol y estableciendo en este texto concreto un paralelismo entre el consumo excesivo de alcohol y el consumo excesivo de agua. Sólo la única participante en el foro de lengua materna alemana se percató de la alusión a la locución alemana, pero solamente después de recibir algunas pistas por parte del profesor. ¡Ojo!: aunque, en rigor, “Sobre la sed” o “Acerca de la sed” deban considerarse infratraducciones, es innegable que <i>funcionan</i> en el texto meta contemplado como un todo y pueden aceptarse si los estándares del traductor o el revisor no son muy elevados. Lo curioso es que el chatbot optara por no proponer ninguna equivalencia para el título. Preguntado por este particular, ChatGPT admitió el error y se disculpó por él —como es su <i>costumbre</i>— y propuso una traducción en línea con la idea de “beber en exceso”, concretamente, “Sobrepasar el límite”.	

Este ejercicio de microcrítica dio paso, en la segunda parte del foro, a una reflexión más amplia y con un enfoque más abierto sobre el lugar que ocupa o debe ocupar la IA en la práctica de la traducción y su enseñanza.

4.2 Los resultados del foro: exploración cualitativa de actitudes y reflexiones sobre la IA en traducción

Si la primera parte del foro se articuló en torno a un ejercicio de microcrítica de la traducción bastante parecido a los que se desarrollan habitualmente en el aula, la segunda parte abrió un espacio de diálogo más libre, en el que

los participantes se alejaron del análisis textual concreto para reflexionar de manera más general y personal sobre la IA en el mundo de la traducción. El presente apartado propone una exploración cualitativa de lo argumentado, con el objetivo de identificar experiencias, actitudes, preocupaciones y expectativas que permitan comprender cómo se está configurando el debate en torno a estas tecnologías tanto en el ámbito docente como en el profesional. Concebida como una conversación guiada, la actividad propició toda una serie de intervenciones que permiten esbozar un panorama matizado de las actitudes de los participantes ante el uso de la IA en traducción. Como en la controversia entre apocalípticos e integrados, lejos de conformar un corpus homogéneo, las posturas expresadas oscilan entre una reticencia abierta y un entusiasmo moderado, pasando por posiciones intermedias de aceptación crítica. Sin aventurar conclusiones, a continuación analizamos una selección de intervenciones atendiendo a siete ejes temáticos recurrentes.

En primer lugar, en relación con (1) el uso efectivo de los chatbots, se observa una clara diversidad de prácticas. Dani, con 28 años el mayor del grupo, asegura no haber utilizado nunca este tipo de herramientas y expresa un rechazo de principio: “Yo no lo he usado nunca. [...] Ni ChatGPT ni nada de eso. Es una cosa que no me atrae. [...] A mí la idea de hablar con una máquina se me hace muy cuesta arriba.” Una actitud similar, aunque menos radical en la práctica, se aprecia en Anna, de 27 años, quien afirma: “Yo intento no utilizarlos. Si utilizo alguno, es ChatGPT, pero es que no me gusta, lo odio, por decirlo así”. No obstante, su experiencia profesional la llevó a emplear ChatGPT por indicación de su empleador, con resultados que valora positivamente en tareas de documentación y terminología en un ámbito especializado —el póquer—, si bien subraya de forma explícita que “ChatGPT nunca vio el texto original”. Elena, de 24 años, adopta una postura distinta: utiliza ChatGPT de manera habitual, aunque no para traducir directamente, sino como herramienta de apoyo a la redacción. Tal como explica, “siempre parto del texto y luego lo paso para que me lo corrija”, insistiendo en que el contenido y las ideas son propias y que la aportación de la IA se limita a la mejora de la expresión. Por su parte, Víctor, con 21 años el participante más joven, es también quien hace un uso más frecuente de los chatbots. Emplea principalmente ChatGPT —y, en menor medida, Copilot— tanto para traducir como para documentarse; y

destaca su utilidad en diversos encargos académicos y también en su reciente trabajo de fin de grado, especialmente en términos de ahorro de tiempo y apoyo a la redacción. Sin dejar de mostrar una actitud crítica, Víctor describe la experiencia de traducir con ChatGPT como “más amena”.

Un segundo eje temático relevante es el relativo a (2) las instrucciones o *prompts* que los participantes consideran necesarios para obtener resultados satisfactorios. Elena señala que, en caso de usar ChatGPT para traducir, “le daría el encargo para que sepa qué se pide de la traducción y el público y demás”. Anna, por su parte, plantea una cuestión de especial interés desde el punto de vista didáctico al preguntarse qué ocurriría si se pidiera al chatbot que marcara los pasajes sobre los que no se siente seguro. Víctor aporta aquí una reflexión especialmente elaborada basada en su experiencia concreta: explica que la calidad de las traducciones mejora notablemente cuando se proporcionan indicaciones contextuales precisas. En su trabajo con los textos literarios de un autor de lengua alemana de finales del siglo XIX no demasiado conocido, la ausencia de este tipo de información llevaba a ChatGPT a atribuir erróneamente los fragmentos a autores netamente canónicos como Thomas Mann o Wilhelm von Humboldt. Para evitarlo, optó por especificar de forma sistemática datos sobre el autor y su época, así como por utilizar instrucciones como: “Traduce según los estándares de traducción literaria o para que suene *natural* en español sin alejarte del sentido”. En su opinión, sin este tipo de *prompts*, ChatGPT tiende a traducir de manera excesivamente literal o a introducir elementos inventados.

Obviamente, la conversación también abordó el tema de (3) la calidad de las traducciones generadas por ChatGPT. Dani se muestra positivamente sorprendido por las traducciones analizadas en la primera parte del foro —aunque, desde su escepticismo, considera que deben de constituir la excepción— y, en general, mantiene su valoración negativa de la traducción automática. Relata su experiencia como lector de textos en español claramente traducidos por la máquina, en los que detecta patrones recurrentes y deficiencias en cuanto a la legibilidad y la inteligibilidad: “No me entero de nada de lo que están diciendo, y me cargo la página en el idioma original y se ve a la legua que está traducida por una máquina y es más difícil la comprensión”. Elena, en cambio, tiene la impresión de que ChatGPT “quizás ya esté un paso por delante” de la

traducción automática *tradicional* y de que “ya es otra cosa que tiene mucha más calidad”. Víctor califica la calidad de las traducciones de ChatGPT como “regular” y considera que, en comparación con las traducciones humanas, “está en el estándar”. A su juicio, el rendimiento de la IA es particularmente adecuado en textos científicos, “porque tienen estructura y al final la mayoría de las palabras son monosémicas”, mientras que identifica mayores dificultades en ámbitos como el jurídico o el literario, debido a la mayor presencia de metáforas o elementos culturales y a la dificultad que puede representar la intencionalidad autoral. Anna, finalmente, es la participante que expresa mayor desconfianza hacia ChatGPT “como traductor” y sostiene que, en muchos casos, un traductor humano tardaría más en corregir una traducción generada por la IA que en traducir el texto desde cero. No obstante, reconoce su utilidad como herramienta de apoyo, especialmente para la documentación, ya que “acelera el proceso”.

Esta cuestión enlaza de manera natural con el debate sobre (4) la eficiencia y el ahorro de tiempo asociados al uso de la IA. Elena considera que ChatGPT “te ahorra mucho tiempo”, tiempo que puede invertirse en otras fases del proceso traductor, y llega a aseverar que “es más eficiente que una persona”. Esta idea conecta con el argumento recurrente clásico en los discursos sobre la automatización, en los que la mayor rapidez de las máquinas se aduce como una de sus principales ventajas frente al trabajo humano.

En el foro se discutieron asimismo (5) las ventajas y riesgos del uso de ChatGPT en la formación de traductores. Elena reconoce que la herramienta “ayuda a entregar el trabajo que tienes que hacer”, pero advierte de que “quizá no ayude a que tú aprendas”. En su opinión, un uso excesivo puede impedir el desarrollo de competencias fundamentales, como la documentación o la evaluación crítica de fuentes, lo que podría derivar en una formación menos sólida. Anna, por su parte, señala efectos negativos sobre la memoria y el esfuerzo cognitivo, al disponer siempre de la IA como apoyo inmediato, y apunta a que puede darse un fenómeno comparable al llamado Efecto Google (amnesia digital o externalización de la memoria). Anna introduce, además, una reflexión de especial interés para los estudios lingüísticos, al observar que el contacto prolongado con textos generados por IA puede influir en la per-

cepción de la *naturalidad* del lenguaje: “Empieza a sonarme *natural* y ya no lo noto tanto como antes, y se está metiendo en nuestro lenguaje”.

Por descontado, otro de los temas tratados fue el de (6) la utilidad didáctica del ejercicio de microcrítica realizado en la primera parte del foro. Dani y Anna se muestran más bien escépticos al respecto. No sin cierta resignación, Dani manifiesta: “A ver, yo creo que hablar de una traducción, de cualquier manera, aporta”. Por su parte, Anna considera que sería más provechoso trabajar con textos producidos por los propios estudiantes: “Para mí, como estudiante, sería mejor coger un texto de cualquiera de la clase y corregirlo juntos”. Para ella, la ventaja principal del enfoque más convencional reside en la posibilidad de debatir directamente con el autor de la traducción (aquí cabría apuntar que ChatGPT sí que justifica sus decisiones si se le pide, algo que no hicimos en el foro por razones logísticas y de planteamiento del ejercicio). Víctor, el más *integrado* de los participantes, adopta una postura más pragmática y subraya la necesidad de incorporar la IA a la enseñanza, dado su uso generalizado: “La inteligencia artificial es una realidad, y se utiliza en clase, y muchísima gente la utiliza, ¿no? Entonces creo que sí que es importante enseñar lo que significa para la divulgación científica y ver hasta dónde es capaz de traducir ciertas cosas o no”.

Por último, el foro puso de relieve (7) una diferencia en la actitud crítica según el origen de la traducción. Dani afirma que aceptaría sin cuestionarlas las traducciones de traductores prestigiosos, pero que se sentiría legitimado para criticar tanto traducciones generadas automáticamente como traducciones de compañeros, desde la confianza en que él podría haberlo hecho mejor. Víctor, Anna y Elena, en cambio, reconocen ser más indulgentes con los errores humanos que con los cometidos por la IA. Como señala Víctor, a los compañeros “les permitimos errores porque son humanos”. Elena interpreta esta mayor severidad hacia la máquina como una forma de reivindicar el valor de la traducción humana: “Al final a la máquina le intentamos sacar errores para decir ‘nosotros somos mejores’”. Anna adopta un tono más combativo al afirmar: “Soy más estricta con la máquina, porque pienso, si ya me sustituye, que por lo menos lo haga bien”. Y concluye con una pregunta retórica reveladora: “¿Para qué utilizarla, si no hace el trabajo a la perfección?”.

En conjunto, los resultados del foro evidencian un panorama de actitudes ante la IA en el ámbito de la traducción que no puede ser más diverso: aunque no contábamos con muchos participantes, en sus manifestaciones coexisten la aceptación pragmática y la resistencia abierta —y abundan las reservas y las cautelas—. Desde el punto de vista metodológico, esta diversidad, lejos de suponer una debilidad, refuerza el valor del foro como espacio de reflexión didáctica. Cabe señalar que la exploración cualitativa de las intervenciones no sólo revela opiniones divergentes, sino también procesos incipientes de negociación de actitudes, expectativas e incluso valores vinculados a la práctica traductora y a su enseñanza. El análisis de los resultados no arroja respuestas concluyentes, pero, en nuestra opinión, sí subraya inequívocamente la conveniencia de integrar la IA en la reflexión didáctica, siempre desde una perspectiva crítica y contextualizada. Partiendo de esta consideración básica, en el apartado que sigue se sintetizan las conclusiones de nuestra investigación y se apuntan ideas que pueden ser de alguna relevancia para la discusión y la innovación docente en traducción.

5 Conclusiones: 12 tesis sobre el uso de la IA en la enseñanza de la traducción

En coherencia con el enfoque cualitativo y exploratorio de este trabajo, para las conclusiones, hemos optado deliberadamente por una estructura algo menos convencional en el ámbito de nuestros estudios y las hemos articulado en forma de tesis emergentes. Estas tesis se derivan del análisis conjunto de los resultados expuestos en los apartados anteriores: no son meras formulaciones tentativas, pero tampoco aspiran a clausurar el debate, sino, por el contrario, a abrirlo aún más, proponiendo líneas de reflexión que fomenten la discusión crítica y constructiva en torno al uso de la IA en la traducción y su enseñanza (y si sirven de punto de partida para futuras investigaciones y experiencias docentes, mucho mejor). Sin pretensión de recoger todo lo visto en nuestra investigación, las tesis permiten recorrer con diligencia la senda marcada por el artículo, que parte de una disquisición muy amplia, de naturaleza casi filosófica, baja luego

a lo más particular y concreto con el ejercicio de microcrítica de la traducción y asciende por último hacia lo más experiencial con las reflexiones personales de los participantes en la segunda parte del foro. Ahí van, pues, nuestras tesis sobre el uso de la IA en la enseñanza de la traducción:

Tesis n.º 1: Los foros de conversación temáticos articulados en torno a la IA y la traducción constituyen un espacio privilegiado para observar procesos de construcción colectiva de conocimiento y para promover una reflexión metadiscursiva sobre la praxis traductora contemporánea y su enseñanza: en los foros *se aprende*.

Tesis n.º 2: El uso de la IA en traducción no parece producir una actitud uniforme entre los estudiantes o jóvenes profesionales, sino todo un abanico de posiciones que van desde el rechazo por principio hasta una adopción moderadamente entusiasta de las nuevas tecnologías, en un *continuum* bastante heterogéneo. Desde nuestro punto de vista, realista pero no exento de escepticismo, tranquiliza constatar que la mayoría de los participantes coincide en su concepción de la IA como herramienta auxiliar, indudablemente valiosa en fases como la documentación, las búsquedas terminológicas o la reformulación estilística pero nunca capaz de sustituir al traductor humano cuando lo que se persigue es un trabajo de calidad.

Tesis n.º 3: La eficacia de la IA en traducción está estrechamente ligada a la capacidad del usuario para formular instrucciones precisas, lo que convierte el diseño de *prompts* en una competencia relevante que debería abordarse explícitamente en la formación: para sacarle el mayor partido a la IA, lo ideal es que el traductor sea capaz de trabajar en diálogo con la máquina.

Tesis n.º 4: El uso continuado —o, si se quiere, abuso— de herramientas basadas en la IA puede tener efectos indeseados sobre la cognición y el lenguaje, esto es, sobre el desarrollo de competencias fundamentales para el traductor: si no se ejercita adecuada y frecuentemente, nuestra capacidad memorística sufre; si se delegan en la IA sistemáticamente las búsquedas léxicas o conceptuales y las reformulaciones, nuestra competencia lingüística mengua; si esta-

mos permanentemente expuestos a los textos redactados por la máquina, acabaremos percibiendo como *natural* el lenguaje que aquella genera —aunque ese lenguaje se aleje del modo de hablar o escribir genuinamente humano—. Desde la didáctica de la traducción, hay que tener muy en cuenta estos riesgos.

Tesis n.º 5: Los ejercicios de microcrítica aplicados a traducciones generadas por la IA resultan didácticamente productivos no tanto por lo que puedan aportar las propias soluciones propuestas por la IA, sino por las discusiones que generan en torno a reproducción del sentido, corrección formal, adecuación estilística ... y también responsabilidad traductora.

Tesis n.º 6: Sabemos que la IA se equivoca al traducir y hemos aportado una cumplida cuenta de gazapos en el apartado de los resultados del ejercicio de microcrítica de la traducción. Ahora bien, el hecho de que los sistemas de traducción basados en IA generen errores y soluciones textuales discutibles no constituye una limitación didáctica, sino un recurso pedagógico de primer orden, en tanto que el análisis de dichos errores favorece el desarrollo de competencias analíticas, evaluativas y metatraductológicas muy indicadas en la formación académica del traductor. Paradójicamente, es la recurrencia de errores en la traducción automática basada en la IA uno de los factores que justifican su integración en la enseñanza de la traducción como instrumento para el entrenamiento del juicio crítico y de las competencias profesionales emergentes.

Tesis n.º 7: Somos conscientes de que puede parecer ingenuo aplicar a la traducción generada por una IA, con sus errores y sus hallazgos, lo que en crítica de la traducción humana puede ser una línea de pensamiento pertinente, incluso reveladora: sabemos que la generación basada en redes neuronales no está sustentada por ninguna clase de razonamiento equiparable al razonamiento humano. No obstante —y con esto se abre un interesante campo de discusión en la didáctica de la traducción—, criticar textos escogidos generados por la IA como si fueran producto de la mente humana sin intermediarios, aunque parezca contraintuitivo, puede ayudar a visibilizar el proceso de traducción

humano (aunque operemos, insistimos, con textos generados por una inteligencia, en rigor, *alienígena*).

Tesis n.º 8: Como pone de manifiesto la traducción del pasaje n.º 5 del apartado 4.1 de este artículo, tabla 5 —y sabemos que sucede a menudo—, todo indica que la IA también tiende a alucinar cuando al traducir se encuentra con datos, como títulos de libros, que no conoce: entonces los genera igual pero lo que nos dice no es cierto, aunque suene plausible. Aquí es donde el traductor tiene que operar con escepticismo y documentarse adecuadamente para proporcionar datos correctos. Lógicamente, para detectar que la IA se ha equivocado, el traductor o estudiante de traducción tiene que saber que la IA *puede equivocarse* —y en *qué*, *cuándo*, *dónde*, *cómo* y hasta *por qué* o incluso *para qué*—.

Tesis n.º 9: Pese a los errores que comete, es una realidad incontrovertible que la IA es capaz de identificar y resolver problemas de traducción complejos, una capacidad que legitimaría su uso en una enseñanza de la traducción que intentara alejarse de la más usual didáctica del error. Suele suceder, sin embargo, que los estudiantes sólo se percatan de la pericia del chatbot al localizar y solucionar problemas con la ayuda del profesor: en efecto, suele ser el profesor quien visibiliza los problemas resueltos por el chatbot y quien aprecia en su justa medida sus aciertos. Por mucho que se trabaje con la IA generativa más afinada, en el aula de traducción, la figura del profesor resulta irremplazable.

Tesis n.º 10: A nadie se le escapa que actualmente, en traducción profesional, se hace un uso extensivo e intensivo de la IA; no obstante, para garantizar la calidad del producto final, sigue siendo necesaria una fuerte capa de posesición humana. El resultado lógico de estas prácticas es que los traductores cada vez se van convirtiendo más en *poseditores*, un oficio en rigor diferente del del traductor primigenio. El riesgo más conspicuo de esta evolución reside en que, si los traductores devenidos poseditores pierden de vista el proceso completo de la traducción en sí y, sobre todo, los textos origen, no es improbable que el sentido de los textos término producidos al alimón por la IA y los traductores/poseditores acabe alejándose del verdadero sentido de los textos

origen... con lo cual quedaría definitivamente postergada la función principal de la traducción. Cotejemos, pues, permanentemente en nuestras aulas el texto origen y sus sucesivas versiones en la lengua meta, sean éstas generadas por la máquina o por nuestros alumnos humanos. Por abundar desde un punto de vista ligeramente diferente en lo antedicho —y expresarlo de paso de manera más simple—, es indudable que al menos parte de las traducciones que genera la IA es subóptima o incorrecta, pero casi nadie se da cuenta porque casi nadie coteja textos origen y textos término (y también porque se le atribuye a la IA una competencia que la IA no tiene —al menos, no todavía—).

Tesis n.º 11: Si lo que se pretende es integrar la IA generativa en la enseñanza de la traducción, desde un posicionamiento docente crítico y humanístico, no se trata en primera instancia de *enseñar al estudiante a usar la IA para traducir*, sino, desde el punto de vista del enseñante, de *usar la IA para enseñar a traducir*. La IA debe ser entendida, pues, como super recurso capaz de generar borradores y prestar toda clase de servicios de asistencia docente, pero no como herramienta destinada a redactar versiones definitivas: ésta es y debe seguir siendo la labor del traductor humano, al menos durante el aprendizaje del oficio.

Tesis n.º 12: Esta última tesis parece de Perogrullo, pero es inexcusable porque se sigue de todas las que anteceden. Nos guste más o nos guste menos, la enseñanza de la traducción no puede cerrarse en banda ante la pujanza de la IA generativa. Nosotros sostenemos que el uso de la IA como herramienta asistencial en la didáctica de la traducción puede contribuir significativamente a una enseñanza más analítica, consciente y orientada a la calidad, pero también subrayamos que el papel del docente humano sigue siendo indispensable: es la orientación del profesor la que permite al estudiante comprender plenamente los retos de la traducción y desarrollar una actitud verdaderamente crítica y reflexiva. En resumen, sí a la IA en el aula de traducción pero sólo cuando se inserta en un marco pedagógico firmemente dirigido por el profesor.

Agradecimientos

Los autores desean expresar su más sincero agradecimiento a D.^a Elena Alonso de la Cruz, D. Daniel Barea Pozo, D.^a Anna Nagel y D. Víctor Onsurbe de la Varga, estudiantes de posgrado y jóvenes traductores e intérpretes que participaron con gran entusiasmo y compromiso en el foro de conversación analizado en este artículo. Su activa implicación, la energía y pertinencia de sus intervenciones y el sobresaliente trabajo de preparación previa fueron fundamentales para la ejecución de nuestro proyecto, que sin su concurso hubiera resultado inviable: vaya para ellas y ellos nuestro más especial reconocimiento.

Bibliografía

- BARBOUR, ROSALINE (2013): *Los grupos de discusión en investigación cualitativa*. Madrid: Ediciones Morata.
- BERNARDINI, SILVIA / GASPARI, FEDERICO (2025): “The Evolution of Translation (Studies) in the Technocene.” *Translation and Translanguaging in Multilingual Contexts* 11: 237–251.
- BOWKER, LYNNE / BUITRAGO CIRO, JAIRO (2019): *Machine Translation and Global Research: Towards Improved Machine Translation Literacy in the Scholarly Community*. Bingley: Emerald Publishing Limited.
- BUSCH-LAUER, INES-ANDREA (2015): “Was einen guten Fachsprachenlehrer auszeichnen sollte ...” *Facetten der Fachsprachenvermittlung Englisch: Hands on ESP Teaching*. Ed. Ines-Andrea Busch-Lauer. Berlin: Frank & Timme. 27–42.
- CALSAMIGLIA, HELENA / VAN DIJK, TEUN A. (2004): “Popularization Discourse and Knowledge about the Genome.” *Discourse & Society* 15.4: 369–389. <http://www.jstor.org/stable/42888625> (8 de enero de 2026).
- COSERIU, EUGENIO (1977): “Lo erróneo y lo acertado en la teoría de la traducción.” *El hombre y su lenguaje. Estudios de teoría y metodología lingüística*. Ed. Eugenio Coseriu. Madrid: Gredos. 214–239.
- DRÖSSER, CHRISTOPH (2004): “Über den Durst. Stimmt es, dass man sterben kann, wenn man zu viel Wasser trinkt?” *Die Zeit* 01/2005. https://www.zeit.de/2005/01/Ueber_den_Durst (11 de diciembre de 2025).
- DRÖSSER, CHRISTOPH (2005): “Schwache Hausmittel.” *Die Zeit* 16/2003. <https://www.zeit.de/2003/16/Stimmtsrotweinflecken> (6 de diciembre de 2025).

- DRÖSSER, CHRISTOPH (2011): "Stirbt man, wenn Luft in die Adern gespritzt wird?" *Die Zeit* 29/2011. <https://www.zeit.de/2011/29/Stimmts-Luft-Spritze> (11 de diciembre de 2025).
- DRÖSSER, CHRISTOPH (2014): "Verhindert der Genuss von Alkohol Magen- und Darminfektionen?" *Die Zeit* 52/2014. <https://www.zeit.de/2014/52/alkohol-in-fektion-stimmts> (11 de diciembre de 2025).
- DRÖSSER, CHRISTOPH (2016): "Kann man einem Toten die Augen schließen, indem man sanft darüberstreicht?" *Die Zeit* 04/2016. <https://www.zeit.de/2016/04/tod-auge-schliessen-fernseh-klischee-stimmts> (11 de diciembre de 2025).
- DRÖSSER, CHRISTOPH (2017): "1.000 Folgen? Stimmt!" *Die Zeit* 02/2017. <https://www.zeit.de/wissen/2017-01/sekt-loeffel-badewanne-erkaeltung-frieren-stimmts-1000> (11 de diciembre de 2025).
- DRÖSSER, CHRISTOPH (2022): "Es reicht, wenn man sich nur mit Wasser wäscht. Stimmt's?" *Die Zeit* 01/2022. <https://www.zeit.de/2022/01/natuerlich-waschen-james-hamblin-seife-shampoo> (11 de diciembre de 2025).
- DRÖSSER, CHRISTOPH (2024): *Was KI mit unserer Sprache macht*. Berlin: Dudenverlag.
- DRÖSSER, CHRISTOPH (2025): "Alkoholiker brauchen höher dosierte Narkosen. Stimmt's?" *Die Zeit* 14/2025. <https://www.zeit.de/2025/14/anaesthesie-narkose-alkoholiker-dosierung-forschung> (6 de diciembre de 2025).
- ECO, UMBERTO (1964): *Apocalittici e integrati: comunicazioni di massa e teorie della cultura di massa*. Milano: Bompiani.
- ECO, UMBERTO (2023): *Apocalípticos e integrados*. Barcelona: Lumen.
- EMT BOARD = EUROPEAN MASTER'S IN TRANSLATION BOARD (2022): *EMT Competence Framework 2022*. https://commission.europa.eu/news-and-media/news/updated-version-emt-competence-framework-now-available-2022-10-21_en (6 de diciembre de 2025).
- GARRISON, DONN RANDY / ANDERSON, TERRY (2010): *El e-learning en el siglo XXI*. Barcelona: Ediciones Octaedro.
- GUTIÉRREZ RODILLA, BERTHA M. (2005): *El lenguaje de las ciencias*. Madrid: Gredos.
- HANSEN-SCHIRRA, SILVIA / KIRALY, DONALD (EDS.) (2013): *Projekte und Projektionen in der Translatorischen Kompetenzentwicklung*. Frankfurt am Main: Lang.
- HATIM, BASIL / MASON, IAN (1995): *Teoría de la traducción: una aproximación al discurso*. Barcelona: Ariel.

- HOUSE, JULIANE (2015): *Translation Quality Assessment: Past and Present*. London: Routledge.
- KELLY, DOROTHY (2005): *A Handbook for Translator Trainers: A Guide to Reflective Practice*. Manchester: St. Jerome.
- KIRALY, DONALD (2012): "Growing a Project-Based Translation Pedagogy: A Fractal Perspective." *Meta* 57.1: 82–95. DOI: 10.7202/1012742ar.
- KIRALY, DONALD (2014): *A Social Constructivist Approach to Translator Education: Empowerment from Theory to Practice*. London: Routledge.
- LAHBY, MOHAMED (ed.) (2024): *General Aspects of Applying Generative AI in Higher Education: Opportunities and Challenges*. Cham: Springer.
- MOORKENS, JOSS / CASTILHO, SHEILA / GASPARI, FEDERICO / DOHERTY, STEPHEN (eds.) (2018): *Translation Quality Assessment: From Principles to Practice*. Cham: Springer.
- MUÑOZ MARTÍN, RICARDO / OLALLA-SOLER, CHRISTIAN (2022): "Translating Is Not (Only) Problem Solving." *The Journal of Specialised Translation* 38: 3–31. DOI: 10.26034/cm.jostrans.2022.081.
- MYERS, GREG (2003): "Discourse Studies of Scientific Popularization: Questioning the Boundaries." *Discourse Studies* 5.2: 265–279. DOI: 10.1177/1461445603005002006.
- PRADAS MACÍAS, ESPERANZA MACARENA / ZWISCHENBERGER, CORNELIA (2021): "Quality and Norms in Conference Interpreting." Eds. Michaela Abl-Mikasa / Elisabet Tiselius. *The Routledge Handbook of Conference Interpreting*. London: Routledge. 243–257.
- RAFNER, JANET / DELLERMANN, DOMINIK / HJORTH, ARTHUR / VERASZTO, DORA / KAMPF, CONSTANCE / MACKAY, WENDY / SHERSON, JACOB (2021): "Deskilling, Upskilling, and Reskilling: a Case for Hybrid Intelligence." *Morals & Machines* 1: 24–39.
- SCHNABEL, ULRICH / SPIEWAK, MARTIN (2022): "Wissensdurst und Neugier sind unerschöpflich." *Die Zeit* 26/2022. <https://www.zeit.de/2022/26/stimmmts-kolumne-fragen-wissenschaft-christoph-droesser-interview> (6 de diciembre de 2025).
- DE SEMIR, VLADIMIR (2015): *Decir la ciencia: divulgación y periodismo científico de Galileo a Twitter*. Barcelona: Universitat de Barcelona.
- SNELL-HORNBY, MARY (1999): *Handbuch Translation*. Tübingen: Stauffenburg.
- VIEIRA, ANA / MESQUITA, ANABEL (2025): "Generative Artificial Intelligence in Higher Education: Challenges, Opportunities and Pedagogical Implications." *Journal of Technologies Information and Communication* 5.1. DOI: 10.55267/rtic/16675.

Moderner Übersetzungsunterricht unter Nutzung digitaler Technologien

Abstract

The rapid development of Artificial Intelligence (AI) has profoundly transformed translation practices and poses new challenges for higher education. This article presents a multi-stage teaching project designed for MA students that critically explores the potentials and limitations of AI-based translation tools. Through the comparative analysis of literary and contemporary texts—most notably Franz Kafka’s short prose piece *“Gib’s auf”*—students examine outputs from tools such as DeepL, Google Translate, ChatGPT, and Matecat. The project combines text analysis, comparison of machine translations, post-editing, and collaborative translation workshops, including direct interaction with a contemporary author. Results show that while AI tools provide efficient and often accurate support, they remain limited in conveying stylistic nuance, pragmatics, and cultural meaning. The study demonstrates that translation quality depends on human interpretation, creativity, and responsibility, highlighting AI as a didactic catalyst rather than a substitute for professional translational competence.

Keywords Artificial Intelligence, literary translation, post-editing, translation didactics, Translation tools

1 Einführung

Die rasante Entwicklung Künstlicher Intelligenz (KI) hat den Bereich des Übersetzens grundlegend verändert und stellt Lehrende wie Lernende vor neue Herausforderungen. Maschinelle Übersetzungsprogramme sind längst

fester Bestandteil des akademischen Alltags und eröffnen vielfältige Möglichkeiten für Sprachunterricht, Textanalyse und kreative Produktionsprozesse. Gerade im Hochschulkontext zeigt sich jedoch, dass der selbstverständliche Einsatz dieser Werkzeuge häufig mit einem Mangel an kritischer Reflexion einhergeht.

Der vorliegende Artikel nimmt diesen Befund zum Anlass, um ein mehrstufiges Unterrichtsprojekt vorzustellen, das Masterstudierenden im Studiengang Germanische Sprachen und Literaturen (einschließlich Translation/Übersetzen) mit der ersten Fremdsprache Deutsch exemplarisch die Potenziale und Grenzen KI-gestützter Übersetzungstools vermittelt. Anhand literarischer und zeitgenössischer Texte, der Analyse maschineller Übersetzungen sowie der eigenen produktiven Arbeit wird sichtbar, wie KI den Übersetzungsprozess unterstützen kann – jedoch niemals die interpretierende, reflektierende und kreativ entscheidende Rolle des Menschen ersetzt. Die Darstellung zeigt damit nicht nur die didaktische Relevanz solcher Projekte, sondern auch, welche neuen Formen des Lernens entstehen, wenn menschliche und maschinelle Textverarbeitung bewusst miteinander verschränkt werden.

Im Unterricht wurde das Thema „Übersetzen mit KI-Tools“ aufgegriffen, um Studierenden einen zeitgemäßen Einblick in moderne sprachliche Arbeitsweisen zu ermöglichen. Zu Beginn stellte die Lehrkraft kurz den Aufbau der Unterrichtseinheit vor und erläuterte, welche Kompetenzen dabei im Mittelpunkt stehen: der reflektierte Einsatz digitaler Werkzeuge, das Vergleichen verschiedener Übersetzungen und das kritische Bewerten von Ergebnissen.

Im praktischen Teil erhielten die Lernenden eine Übersetzungsaufgabe, die sie mithilfe gängiger KI-Tools wie DeepL, Google Translate oder ChatGPT bearbeiten sollten. Ziel war es nicht nur, die schnell und technisch korrekt wirkenden Ergebnisse zu nutzen, sondern sie miteinander zu vergleichen und auf typische Schwächen der Tools bzw. auf die stilistischen Besonderheiten in den Übersetzungen zu achten. Die Lernenden bemerkten dabei schnell, dass Übersetzungen je nach Tool unterschiedlich ausfallen können und dass KI zwar hilfreiche Unterstützung bietet, aber keineswegs unfehlbar ist. Besonders spannend war die Diskussion darüber, wie sehr KI-Übersetzungen den eigenen Schreibstil beeinflussen oder welche Verantwortung beim menschlichen Nachbearbeiten (*Post-Editing*) bleibt.

Zum Abschluss wurde ein Ausblick auf weitere Einsatzmöglichkeiten gegeben. KI-gestützte Sprachtools können Lernende nicht nur bei der Übersetzung unterstützen, sondern auch beim eigenständigen Verfassen von Texten, bei der gezielten Überarbeitung von Formulierungen sowie beim Üben und Vertiefen fremdsprachlicher Kompetenzen, etwa durch Feedback, Korrekturvorschläge und interaktive Dialoge (vgl. UNESCO 2023). Gleichzeitig wurde betont, wie wichtig es ist, diese Technologien bewusst, kritisch und kreativ einzusetzen. Die Unterrichtseinheit zeigte damit eindrucksvoll, wie KI den Spracherwerb bereichern kann, ohne die Bedeutung menschlicher Sprachkompetenz zu schmälern.

2 Übersetzen und KI-Tools

Im Jahr 2023 zeigte eine Umfrage unter kirgisischen Studierenden, wie stark KI-Tools bereits in den Lernalltag eingezogen sind (vgl. Lange/Wolf 2023). Die Ergebnisse machten deutlich, dass die Nachfrage nach digitalen Übersetzungs- und Schreibwerkzeugen enorm hoch ist. Viele Studierende greifen nahezu selbstverständlich auf Anwendungen wie DeepL, Google Translate oder ChatGPT zurück, um Texte zu verstehen, zu übersetzen oder schriftliche Aufgaben schneller zu bewältigen.

Gleichzeitig wurde jedoch sichtbar, dass der Umgang mit diesen Technologien häufig noch recht naiv erfolgt. Die kritische Reflexion, etwa über die Qualität der Übersetzungen, mögliche Fehlerquellen oder Fragen des Datenschutzes, spielt in vielen Fällen kaum eine Rolle. Die Umfrage zeigte zudem, dass Lehrkräfte oft nur über begrenzte Kompetenzen im Bereich KI-gestützten Lernens verfügen. Viele fühlen sich unsicher im Umgang mit den Tools oder sind nicht ausreichend darüber informiert, wie diese sinnvoll und verantwortungsvoll in den Unterricht integriert werden können.

Daraus ergibt sich ein klarer Bedarf nach Fortbildungen – sowohl für Studierende, um einen reflektierten und kompetenten Umgang mit KI zu lernen, als auch für Lehrkräfte, um diese Entwicklungen pädagogisch zu begleiten. Die Ergebnisse unterstreichen, dass KI im Bildungsbereich nicht nur neue

Möglichkeiten eröffnet, sondern auch neue Anforderungen an die digitale und kritische Kompetenz aller Beteiligten stellt.

Im Bereich der automatischen Übersetzung stehen heute zahlreiche Hilfsmittel zur Verfügung: Neben kostenpflichtigen und kostenlosen Varianten von ChatGPT nutzen viele Lernende und Lehrende auch Google Translate, DeepL oder professionelle Tools wie Matecat. Diese Technologien bieten beachtliche Unterstützung im Alltag und erleichtern das schnelle Erfassen fremdsprachiger Inhalte. Dennoch bleiben sie – zumindest derzeit – klar hinter der Leistungsfähigkeit menschlicher Übersetzerinnen und Übersetzer zurück. „ChatGPT and DeepL are useful but imperfect tools for translating literature, and they require human review and improvement“ (Karabayeva/Kalizhanova 2024: 1).

Gerade in Fachkreisen wird immer wieder betont, dass KI keineswegs eine vollwertige Alternative zum professionellen Übersetzen darstellt. Der Bundesverband der Dolmetscher und Übersetzer (BDÜ 2023) warnte bereits 2023 davor, die „fundierte Expertise und Beratungskompetenz“ menschlicher Fachkräfte zu unterschätzen. Maschinelle Sprachübertragung kann zwar große Datenmengen in Sekunden verarbeiten, bleibt aber auf Mustererkennung beschränkt und verfügt über kein echtes Textverständnis. Ein Computerprogramm „versteht“ nicht, was eine bestimmte Textsequenz meint, welche Nuancen entscheidend sind oder wie Kontext und kulturelle Bezüge die Bedeutung verändern. In einer aktuellen Studie aus dem Jahr 2025 mit dem Sprachenpaar Englisch–Türkisch wurden Kurzgeschichten übersetzt. Es wurden 381 Sätze aus Kurzgeschichten maschinell (DeepL) und manuell übersetzt und verglichen. Das Ergebnis waren häufige Unterschiede bei Lexik, Semantik, Grammatik und Stil, insbesondere aber Probleme bei kulturellen Nuancen und Kontext-Abwägungen (Dolmaci 2025: 93).

Auch im Sprachenpaar Kasachisch und Englisch konnte festgestellt werden, dass Künstliche Intelligenz „provides a word-to-word translation, and comparing the earth to the velvet may not fit the English picture of the world“ (Manapbayeva et al. 2024: 527). Die KI konnte im Gegensatz zum menschlichen Übersetzer feine Nuancen im Text nicht ausdrücken (vgl. ebd). Auch Elkins (2024: 6) spricht die fehlende Emotionalität, jedoch auch die große Hilfe von KI bei der Übersetzung an. „More research is needed, but this study suggests that AI can help surface elements of a translation that may not be

visible using a traditional close reading method“ (ebd.). Elkins spricht von Charakteristiken, die „verloren“ gehen, wenn eine KI übersetzt. „While one can never truly replicate the way it feels to read a literary work in the original, both of these methods reveal subtle aspects of the reading experience that can easily be lost in translation“ (Elkins 2024: 6).

Daher gilt weiterhin: KI kann wertvolle Unterstützung leisten, doch hochwertige Übersetzung ist ohne den denkenden, interpretierenden Menschen nicht möglich. Die Technik eröffnet neue Möglichkeiten – ersetzt aber nicht das menschliche Urteilsvermögen, das Sprachgefühl und die inhaltliche Verantwortung, die professionelle Übersetzungen auszeichnen.

3 Kurze Vorstellung des Übersetzungsunterrichts

Für die automatisierte Übersetzung eines Textes stehen verschiedene Werkzeuge zur Verfügung, die je nach Bedarf und Situation eingesetzt werden können. Professionelle Übersetzer greifen häufig auf Programme wie Trados SDL zurück, die umfangreiche Funktionen zur Verwaltung von Terminologie und zu Übersetzungsspeichern bieten. Für Rohübersetzungen von Texten wird oft Google Translate genutzt, das zahlreiche Sprachen unterstützt und leicht zugänglich ist. DeepL ist für seine besonders präzisen und natürlich klingenden Übersetzungen bekannt. Zudem können moderne KI-gestützte Systeme wie ChatGPT eingesetzt werden, die nicht nur übersetzen, sondern auch stilistische Anpassungen oder erklärende Umschreibungen liefern. Diese Vielfalt an Werkzeugen ermöglicht es, je nach Textart und Zielsetzung die passende Übersetzungsmethode auszuwählen.

Im November 2024 wurde ein Online-Unterrichtsmodul mit Masterstudierenden des KPI¹ in Kyjiw durchgeführt. Insgesamt drei Studierende haben während des Unterrichts durch einen geteilten Bildschirm Anfragen an die jeweiligen Chatbots gegeben (im Falle der Schwierigkeiten) und die Ergeb-

.....

1 Nationale Technische Universität „Ihor Sikorsky Kyjiwer Polytechnisches Institut“

nisse wurden im Kurs sofort diskutiert und eingeordnet. Die Seminareinheit umfasste zwei Sitzungen à 80 Minuten und konzentrierte sich auf die Übersetzungsrelevante Textanalyse sowie auf den Vergleich verschiedener Übersetzungen eines literarischen Textes. Als Grundlage diente Franz Kafkas kurze, aber vielschichtige Erzählung „Gib’s auf“, die aufgrund ihrer Dichte, Offenheit und interpretatorischen Rätselhaftigkeit besonders geeignet ist, Unterschiede zwischen dem Übersetzungsprozess des Menschen und eines Programmes sichtbar zu machen.

Die Studierenden setzten sich zunächst mit den zentralen Merkmalen des Ausgangstextes auseinander: stilistischen Feinheiten, semantischen Mehrdeutigkeiten, narrativer Situation und der für Kafka typischen Atmosphäre. Anschließend wurden fünf verschiedene KI-Tools herangezogen, um Übersetzungen zu erstellen und gegenüberzustellen. Dabei kamen sowohl bekannte Anwendungen wie Google Translate, DeepL und ChatGPT als auch spezialisiertere Tools wie Matecat zum Einsatz.

Im Vergleich der Ergebnisse zeigte sich schnell, wie unterschiedlich die Systeme mit Kafkas dichter Sprache umgehen: Manche Tools lieferten überraschend flüssige, aber interpretatorisch stark glättende Übersetzungen, andere produzierten präzise, jedoch stilistisch unpassende Formulierungen. Für die Studierenden bot dieser direkte Vergleich einen aufschlussreichen Einblick in die Möglichkeiten und Grenzen maschineller Übersetzung – und zugleich eine Vertiefung ihres Verständnisses für professionelle Übersetzungsarbeit.

4 Übersetzung des Textes „Gib’s auf“ von Franz Kafka

Im Folgenden soll der Originaltext von Franz Kafkas „Gib’s auf“ gegeben werden, damit anschließend die KI-generierten Übersetzungen mit dem Originaltext verglichen werden können.

Es war sehr früh am Morgen, die Straßen rein und leer, ich ging zum Bahnhof. Als ich eine Turmuhr mit meiner Uhr verglich, sah ich, dass es

schon viel später war, als ich geglaubt hatte, ich musste mich sehr beeilen, der Schrecken über diese Entdeckung ließ mich im Weg unsicher werden, ich kannte mich in dieser Stadt noch nicht sehr gut aus, glücklicherweise war ein Schutzmann in der Nähe, ich lief zu ihm und fragte ihn atemlos nach dem Weg. Er lächelte und sagte: „Von mir willst du den Weg erfahren?“ „Ja“, sagte ich, „da ich ihn selbst nicht finden kann.“ „Gib’s auf, gib’s auf“, sagte er und wandte sich mit einem großen Schwunge ab, so wie Leute, die mit ihrem Lachen allein sein wollen.

Übersetzungsprozess des Textes „Gib’s auf“ von Franz Kafka im Seminar

Nach der Lektüre des Textes wurde den Studierenden die Aufgabe gegeben, zu erraten, wer den Text geschrieben hat. Wenn man nur nach Stil und Sprache urteilt, passt der Text weder zu Goethe, Schiller, Brecht oder Kafka. Wenn man aber einen der vier Autoren auswählen muss, wäre die einzige nachvollziehbare Wahl: Kafka – nicht wegen des Stils, sondern wegen des Inhalts, da der Text sich direkt auf Kafkas Erzählung „Gib’s auf“ bezieht.

Die Studierenden hatten danach die folgenden Aufgaben zu bearbeiten:

- Was haben Sie beim ersten Lesen spontan verstanden?
- Aus welcher Zeit stammt dieser Text (kontextuelle Kriterien der Textanalyse)?
- Welche Art von Text/Textsorte ist das? Konkretisieren Sie.
- Welche Figuren spielen eine Rolle? Beschreiben Sie die Figuren! Was erfahren Sie aus dem Text?
- Welche Wörter sind „Schlüsselwörter“ im Text?

Im Folgenden wird der Text mit Schlüsselwörtern (durch Fettdruck gekennzeichnet) wiedergegeben:

*Es war sehr früh am Morgen, die Straßen rein und leer, ich ging zum Bahnhof. Als ich eine **Turmuh**r mit meiner Uhr verglich, sah ich, dass es schon viel*

*später war, als ich geglaubt hatte, ich musste mich sehr **beeilen**, der **Schrecken über diese Entdeckung** ließ mich ganz unsicher werden, ich kannte mich in dieser Stadt noch nicht sehr gut aus, glücklicherweise war ein Schutzmann in der Nähe, ich lief zu ihm und fragte ihn **atemlos** nach dem Weg. Er lächelte und sagte: „Von mir willst du den Weg erfahren?“ „Ja“, sagte ich, „da ich ihn selbst nicht finden kann.“ „Gibs auf, gib auf“, sagte er und wandte sich mit einem großen Schwunge ab, so wie Leute, die mit ihrem Lachen allein sein wollen.*

Die Studierenden sollten zunächst die Schlüsselwörter herausuchen und im nächsten Schritt den Text in Bezug auf folgende Kriterien untersuchen:

1. Sprachliche Mittel
Wortwahl, Stilfiguren und rhetorische Mittel
2. Satzbau
3. Erzählperspektive, Zeitgestaltung (Zeitdehnung/Zeitraffung), Sprachstil, Tempo und Rhythmus
4. Aufbau des Textes, formale Makrostruktur (Gliederungssignale), Kohärenz und Textstruktur (Absatzstruktur, Erzähltechnik, Zeitsprünge)
5. Intertextualität
6. Emotionale Wirkung
7. Ziel des Textes (Textfunktion)

Der analysierte Text mit Hervorhebungen sieht wie folgt aus:

*Es war sehr früh am Morgen, die Straßen **rein und leer** [Ellipse], ich ging zum Bahnhof.*

*Als ich eine **Turmuhr mit meiner Uhr** [Vergleich] verglich, sah ich, dass es schon viel **später war, als ich geglaubt hatte**, ich musste mich sehr **beeilen**, der **Schrecken über diese Entdeckung ließ mich im Weg** [Inversion] unsicher werden, ich kannte mich in dieser Stadt noch nicht sehr **gut** aus, glücklicherweise war ein Schutzmann in der Nähe, ich lief zu ihm und fragte ihn **atemlos** nach dem Weg. Er lächelte und sagte: „Von*

mir willst du den Weg erfahren?“ „Ja“, sagte ich, „da ich ihn selbst nicht finden kann.“ „Gibs auf, gibbs auf“, sagte er und wandte sich mit einem großen Schwunge ab, so wie Leute, die mit ihrem Lachen allein sein wollen. [Hyperbel]

Nach der Textanalyse, der Markierung unbekannter Wörter sowie von Schlüsselbegriffen gilt es, im Übersetzungsprozess auch die Satzsyntax und die verwendeten stilistischen Mittel zu analysieren. Auch die emotionale Wirkung des Originaltextes sollte analysiert und, soweit möglich, in der Übersetzung bewahrt werden. Eine gründliche Textrevision schließt den Übersetzungsprozess ab.

Im Seminar wurden zunächst folgende Schlüsselwörter von den Studierenden identifiziert:

eine Turmuhr, mich sehr beeilen, der Schrecken über diese Entdeckung, ein Schutzmann, atemlos, den Weg erfahren, ‚Gibs auf, gibbs auf‘, mit einem großen Schwunge. Zunächst werden im Folgenden die fünf Übersetzungen gegeben, danach werden diese hinsichtlich ihrer Qualität bewertet.

Die mit den Programmen durchgeführten Übersetzungen lauten wie folgt:

(1) ChatGPT kostenpflichtige Version

Було дуже рано вранці, вулиці були чисті й порожні, я йшов до вокзалу. Коли я порівняв годинник на вежі зі своїм, то побачив, що вже набагато пізніше, ніж я думав. Мені потрібно було дуже поспішати, і цей страх через відкриття змусив мене розгубитися в дорозі, адже я ще не дуже добре орієнтувався в цьому місті. На щастя, неподалік був поліцейський. Я підбіг до нього і, ледве переводячи подих, запитав у нього дорогу. Він посміхнувся і сказав: «Ти хочеш, щоб я підказав тобі дорогу?» «Так», — відповів я, — «бо сам знайти її не можу». «Залиш це, залиш», — сказав він і відвернувся з широким жестом, як люди, які хочуть посміятися на самоті.

(2) Chat GPT kostenlos

Був дуже ранній ранок, вулиці були чистими та безлюдними, я йшов до вокзалу. Порівнявши час на вуличному годиннику зі своїм, я зрозумів, що вже набагато пізніше, ніж я думав. Я мав дуже поспішати, і це відкриття так налякало мене, що я став невпевненим у своєму маршруті. Я ще погано орієнтувався в цьому місті. На щастя, поруч був поліцейський. Я підбіг до нього і, задихаючись, запитав дорогу. Він усміхнувся і сказав: «Ти хочеш дізнатися дорогу в мене?» «Так,» відповів я, «бо сам не можу її знайти.» «Залиш це, залиш,» сказав він і різко відвернувся, як це роблять люди, які хочуть сміятися на самоті.

(3) Google Translate

Був дуже ранній ранок, вулиці чисті та порожні, я поїхав на вокзал. Коли я порівнював багаторазовий годинник зі своїм, я побачив, що він набагато пізніший, ніж я думав, мені потрібно було поспішити, шок від цього відкриття змусив мене почуватися невпевнено в дорозі, я не знав, як обійти це місто дуже добре, на щастя, поруч був поліцейський, я підбіг до нього і задихаючись запитав у нього дорогу. Він посміхнувся і сказав: «Ти хочеш дізнатися дорогу від мене?» «Так, – сказав я, – оскільки я не можу знайти його». великі гойдалки, як це роблять люди, які хочуть побути наодинці зі своїм сміхом.

(4) Matecat

Було дуже рано вранці, вулиці були чисті і порожні, я поїхав на вокзал. Коли я порівнював вежевий годинник зі своїм годинником, то побачив, що вже набагато пізніше, ніж я вірив, мені доводилося дуже швидко поспішати, жак цього відкриття робив мене невпевненим на шляху, я не дуже добре знав себе в цьому місті, на щастя поруч був охоронець, я підбіг до нього і затамував подих попросив його про дорогу. Він посміхнувся і сказав: «Хочеш знати дорогу від мене?» – Так, – сказав я, – оскільки я не можу знайти його сам, – здавайся, здавайся, – сказав він, відвернувшись з великим розмахом, як люди, які хочуть залишитися наодинці зі своїм сміхом.

(5) DeepL.com (kostenlose Version)

Було дуже рано вранці, вулиці були чисті і порожні, я пішов на залізничний вокзал. Коли я порівняв годинник на вежі зі своїм годинником, то зрозумів, що вже набагато пізніше, ніж я думав, мені треба було поспішати, жаж від цього відкриття змусив мене почуватися невпевнено на своєму шляху, я ще не дуже добре знав це місто, на щастя, неподалік був охоронець, я підбіг до нього і, затамувавши подих, запитав у нього напрямом. Він посміхнувся і сказав: «Хочеш, щоб я показав тобі дорогу?» «Так», – відповів я, – «бо сам не можу її знайти» «Та годі, годі», – сказав він і відвернувся з великим розмахом, як люди, які хочуть залишитися наодинці зі своїм сміхом.

Die fünf Übersetzungen des Textes zeigen deutliche qualitative Unterschiede, sowohl hinsichtlich der semantischen Genauigkeit als auch in Bezug auf Stil, Idiomatik und die Wiedergabe literarischer Nuancen. Am überzeugendsten ist die Übersetzung der kostenpflichtigen ChatGPT-Version. Sie überträgt nicht nur die Bedeutung präzise, sondern spiegelt auch die atmosphärische Dichte und die feinen stilistischen Eigenheiten des Originals wider. Der erzählerische Rhythmus bleibt erhalten, und selbst subtile Wirkungen – etwa der leicht ironische Ton des Dialogs – werden sensibel nachgebildet. Die kostenlose ChatGPT-Version folgt semantisch zwar weitgehend korrekt, zeigt jedoch eine Tendenz zur stilistischen Vereinfachung. Feinheiten wie der bildhafte Ausdruck „mit einem großen Schwunge“ werden weniger literarisch aufgegriffen, wodurch etwas von der charakteristischen Stimmung verloren geht.

DeepL liefert eine insgesamt sehr solide, grammatikalisch saubere Übersetzung, die in vielen Passagen die Bedeutung zuverlässig wiedergibt. Allerdings zeigen sich einzelne kontextuelle Schwächen, etwa bei der falschen Übertragung des Begriffs „Schutzmann“ zu „охоронець“, was im Ukrainischen eine andere Rolle bezeichnet. Diese Art von kultur- und kontextabhängigen Begriffen stellt für das System eine Herausforderung dar, obwohl der übrige Text flüssig und konsistent wirkt.

Deutlich schwächer sind die Übersetzungen von Matecat und Google Translate. Beide wirken stilistisch maschineller und weniger idiomatisch; zudem treten mehrere semantische Fehlzuordnungen auf. Google Translate

erzeugt eine Reihe syntaktisch holpriger oder unpassender Formulierungen, etwa die Fehlübersetzung „великі зойдалки“ für „großen Schwunge“, was die Szene verfremdet. Auch Matecat übersetzt stellenweise zu wörtlich, wodurch der Text unnatürlich wirkt; insbesondere die Wahl des Wortes „охоронець“ statt „полицейський“ führt zu einer Bedeutungsverzerrung. Darüber hinaus gelingt es beiden Systemen weniger, den literarischen Gestus des Originals beizubehalten, sodass der narrative Fluss und die psychologische Dichte der Szene abgeschwächt werden.

Insgesamt zeigt der Vergleich, dass Übersetzungssysteme mit stärker ausgeprägtem kontextuellem Sprachverständnis – wie die kostenpflichtige ChatGPT-Version – deutlich bessere Ergebnisse erzielen. Modelle wie DeepL liefern ebenfalls hohe Qualität, stoßen jedoch bei kontextsensitiven Elementen an Grenzen. Reine maschinelle Übersetzungswerkzeuge ohne tieferes Textverständnis tendieren hingegen zu wörtlichen und syntaktisch mechanischen Lösungen, die dem literarischen Charakter des Ausgangstextes nur bedingt gerecht werden.

5 Übersetzung eines zeitgenössischen Textes mithilfe von KI

Nach der kritischen Analyse maschineller Übersetzungen literarischer Texte (am Beispiel Kafkas) wurde im nächsten Übersetzungsprojekt für dieselben drei Master-Studierenden der Fokus von der rezeptiven zur produktiven Arbeit verlagert. Das zweite Projekt war als Anschluss an das erste vorgesehen, aber diesmal mit dem Fokus auf das Produktive/Kreative und nicht nur eine Analyse von vorgefertigten KI-Übersetzungen. Im zweiten Projekt war KI nur eine „Vorlage“ für das Treffen eigener Übersetzungsentscheidungen. Während ChatGPT zuvor primär als Untersuchungsobjekt diente, wurde es nun in den eigentlichen Übersetzungsprozess integriert – als experimentelles Werkzeug innerhalb eines literarischen Workshops.

Dieser Perspektivwechsel ist didaktisch und theoretisch motiviert: Während die Analyse fremder KI-Übersetzungen vor allem zur Sensibilisierung für

stilistische Defizite, semantische Ungenauigkeiten und Verlust von Pragmatik führte, stand nun die Frage im Zentrum, wie sich KI in den kreativen Akt literarischen Übersetzens einfügt, ohne diesen zu ersetzen. ChatGPT² wurde als Werkzeug eingesetzt, dessen Vorschläge ein Sprungbrett für eigene Entscheidungen bildete.

Besonders relevant ist in diesem Kontext der Wechsel von einem kanonisierten Autor (Kafka) zu einem zeitgenössischen Schriftsteller (Gabriel Wolkenfeld). Durch die direkte Kommunikation mit dem Autor während des Projektes wurde die Autorschaft in ihrer lebendigen Dimension erfahrbar – und damit auch die Grenze zwischen Text, Intention und Übersetzung neu sichtbar.

ChatGPT konnte hier keine autoritative Deutung liefern, sondern lediglich Hypothesen generieren, die sich im Dialog mit dem Originalautor relativier(t)en oder bestätig(t)en. Damit eröffnete sich ein didaktisch produktiver Spannungsraum zwischen menschlicher Interpretation, maschineller Textgenerierung, autoritativem Kontextwissen des Autors und gemeinsamer Textaushandlung im Team. Der Einsatz von ChatGPT als Mustertool wird als Teil einer didaktischen Strategie verstanden, die darauf abzielt, Übersetzungsentscheidungen nicht zu delegieren, sondern genauer zu begründen, zu reflektieren und zu verantworten.

5.1 Kurzzusammenfassung des Textes

Beim zu übersetzenden Text handelt es sich um „Das Ende der Schönheit“ von Gabriel Wolkenfeld. Der Text erzählt von einer Königin, die früher für ihre außergewöhnliche Schönheit verehrt wurde, nun aber verzweifelt versucht, ihr verblassendes Aussehen wiederzuerlangen. Besessen von ihrem Wunsch, erneut die Schönste zu sein, probiert sie unzählige Mittel aus, konsultiert Experten und lässt magische und exotische Materialien zu Kleidern verarbeiten. Doch im Spiegel sieht sie nur ihre eigenen Makel, was sie zunehmend verbittert.

.....
2 Kostenlose Version, GPT-5.1.

Als ihr Selbstbild weiter zerbricht und das Volk beginnt, über ihr Aussehen zu tuscheln, erlässt sie ein Gesetz, das Schönheit verbietet. Mädchen werden verunstaltet, Engelwesen gejagt und Schönheit in versteckte Hinterhöfe gedrängt.

Währenddessen bereitet das Königspaar ein großes Fest für den Besuch des Großkatans und seiner legendär schönen Frau Joffifa vor. Die Königin hofft, Joffifa das Geheimnis ihrer Schönheit zu entlocken. Doch als sie Joffifa schließlich ohne Schleier sieht, erkennt sie, dass auch diese Frau nicht vollkommen ist. Auf die Frage nach ihrem Geheimnis erhält sie lediglich den rätselhaften Satz: *„Ich bin schöner, als ich aussehe.“*

Diese Erkenntnis stürzt die Königin in neue Verzweiflung. Sie begreift weder den Sinn der Aussage noch findet sie darin Trost. Ihr innerer Konflikt bleibt ungelöst, während ihr Hofstaat sie weiterhin mit hohler Schmeichelei umgibt.

5.2 Ziel des Projektes

Ziel des vorliegenden Projektes war die Erstellung einer hochwertigen literarischen Übersetzung ins Ukrainische unter Berücksichtigung der Textsortenspezifika sowie der pragmatischen Funktionen des Textes. Die getroffenen Übersetzungsentscheidungen wurden nach dem funktional-pragmatischen Ansatz von Nord (2011) getroffen. Zentrale Aufgaben umfassten:

- die Entwicklung von Kompetenzen zur übersetzungsrelevanten Textanalyse;
- die Erarbeitung von Strategien zur Übertragung von Anspielungen, impliziten Bedeutungen sowie lexikalischen Einheiten mit geringer Repräsentanz in Lehrwerken;
- die Vereinheitlichung individueller Übersetzungsentscheidungen durch kollektive Redaktionsarbeit.

Im Projekt ließen sich die folgenden Arbeitsphasen abgrenzen:

- Wechselseitige Korrektur der Übersetzungen (Peer-Editing)

- Kollektive Entscheidungen zur Übersetzung von Eigennamen, einzelnen Lexemen sowie zur stilistischen Gesamtgestaltung des Textes
- Kommunikation mit dem Autor: direkte Klärung inhaltlicher und kontextueller Fragen in einer Online-Begegnung mit dem Autor.

5.3 Übersetzungsschwierigkeiten

Im Folgenden werden einige der Übersetzungsschwierigkeiten dargestellt, bei deren Bearbeitung punktuell auf ChatGPT als unterstützendes Instrument zurückgegriffen wurde. Die Tabellen ermöglichen den Vergleich von ChatGPT mit den studentischen Versionen.

5.3.1 Eigennamen

Die Übersetzung von Eigennamen stellt im literarischen Kontext eine besondere Herausforderung dar, insbesondere, wenn Namen semantisch aufgeladen oder symbolisch motiviert sind. Eigennamen in literarischen Texten werden in der Regel nicht übersetzt, sondern in der Originalform belassen.

Tab. 1: Beispiel für die Wiedergabe von Eigennamen

Originaltext	ChatGPT-Übersetzung	Übersetzung von Studierenden
<i>Lukkaschekke</i>	Луккашекке	Луккашекке
<i>Joffifa</i>	Джоффіфа	Йоффіфа

Die studentische Übersetzung dieses Namens ist aus orthografischer und phonetischer Sicht überzeugend. Die Geminata *kk* wurde korrekt durch *кк* wiedergegeben, wodurch die markierte, „harte“ Klangstruktur des Originals erhalten bleibt. Auch die Lautfolge *-schekke* wurde kohärent in *-шекке* übertragen, was der deutschen Phonetik nahekommt. Bei dem Eigennamen *Joffifa* ist die studentische Lösung phonetisch nahe am deutschen *Jo-* (*[jo]*) orientiert und bewahrt die expressive Doppelform *ff* durch *фф*. Die Wahl von *й* betont die Annäherung an die deutsche Aussprache, nicht an eine englische oder ro-

manische Lesart. Die ChatGPT-Übersetzung folgt einer eher englisch beeinflussten Lesestrategie (vgl. „Jack“, „Joe“), die im Deutschen nicht begründet ist.

5.3.2 **Phraseologismen und Idiome**

Idiome und feste Wendungen sind stark kulturgebunden und entziehen sich oft einer wörtlichen Übersetzung. Im Vergleich verdeutlichen die Vorschläge von ChatGPT eine Tendenz zu semantischer Glättung, während die Studierenden stärker auf kreative Neufassungen setzen, um den pragmatischen Effekt im Zieltext zu bewahren.

Tab. 2: Idiomatische Wendungen

Originaltext	ChatGPT-Übersetzung	Übersetzung von Studierenden
<i>das schlägt aufs Gemüt</i>	це важко тисне на душу / це гнітить	серце обвиває гнітюче відчуття
<i>hatte Kopf und Kragen dafür riskiert</i>	ставив на карту власне життя / ризикував головою	ледь не наклав власною головою

Die studentische Version des Idioms *das schlägt aufs Gemüt* führt eine neue metaphorische Struktur ein (*серце обвиває* / *das Herz umwickelt*), die im Original nicht vorhanden ist. Dadurch wird die Aussage emotional intensiviert. Während das Verb *гнітити* im Ukrainischen einen psychischen Zustand knapp und funktional benennt, wird in der studentischen Fassung aus dem verbalen Prädikat ein nominalisiertes Empfinden: *гнітюче відчуття*. Diese Umformung geht über eine bloße semantische Entsprechung hinaus und führt zu einer konzeptuellen Reinterpretation der Aussage. Durch den Wechsel von einem Zustand (*гнітити*) zu einer emotionalen Entität (*гнітюче відчуття*) wird das Leiden nicht mehr als diffuse innere Belastung dargestellt, sondern als ein objektivierbares Gefühl, das das Subjekt „umwickelt“ (*обвиває*). Damit verschiebt sich der Fokus von einem unbestimmten innerpsychischen Druck hin zu einem bildlich verdichteten, poetisch konkretisierten Empfinden. Im Ergebnis ist festzuhalten, dass die studentische Übersetzung weniger reprodu-

ziert als vielmehr literarisch interpretiert, was als kreativer Eingriff gewertet werden kann.

5.3.3 Umgangssprache, vulgäre Lexik

Diese Lexik erfordert besondere Aufmerksamkeit für die semantische Präzision und stilistische Differenzierung, da die unten verwendeten Lexeme stark konnotiert sind und unterschiedliche Nuancen von sozialer Bewertung transportieren. Eine sorgfältige Übersetzung in die Zielsprache muss sowohl die lexikalische Vielfalt als auch die rhetorische Wirkung berücksichtigen.

Tab. 3: Beispiel für Umgangssprache

Originaltext	ChatGPT-Übersetzung	Übersetzung von Studierenden
<i>Grobiane, Gesindel</i>	Грубіяни, голота	неотесані, чернь
<i>Pöbel</i>	чернь	чернь

Die zu übersetzenden deutschen Lexeme *Grobiane*, *Gesindel*, *Pöbel* gehören zu einem abwertenden, soziolinguistisch markierten Wortfeld. Sie variieren in Nuancen zwischen Verhaltensbeschreibung (*Grobiane*), sozialer Abwertung (*Gesindel*) und kollektiver Pejoration (*Pöbel*). Eine gute Übersetzung sollte diese Abstufung, Konnotation und soziale Derbheit adäquat abbilden. In der maschinellen Übersetzung trifft *грубіяни* den Sinn von *Grobiane* sehr gut: Menschen mit grobem, ungehobeltem Verhalten. *Голота* ist ein abwertender Begriff für *Gesindel*, sozial tiefstehende oder verachtete Gruppen, ist semantisch passend und historisch tief verankert im Ukrainischen. *Чернь* ist eine klassische Übersetzung für *Pöbel* und trifft die kollektive, massenhafte Abwertung exakt. Die KI-Version wirkt lexikalisch präzise, stilistisch abgewogen und zielsprachlich natürlich. Sie transportiert sowohl die soziale als auch die emotionale Konnotation des deutschen Originals.

In der studentischen Übersetzung *неотесані* bedeutet wörtlich „unbehauen/ungehobelt“ und entspricht dem Lexem *Grobiane*. *Чернь* wird zweimal verwendet (für *Gesindel* und *Pöbel*), was teilweise zur Bedeutungsnivellie-

rung führt: Zwei deutsche Begriffe unterschiedlicher Konnotation werden im Ukrainischen zu einem einzigen verschmolzen.

Die maschinelle Übersetzung war für Studierende didaktisch wertvoll, weil sie lexikalische, stilistische und semantische Differenzierung sichtbar macht. Die maschinelle Version zeigt, wie man Abstufungen erhält und drei verschiedene Bedeutungsfelder zielsprachlich differenziert wiedergeben kann.

5.3.4 Märchenwelt/Märchenlogik, Absurdität

Diese Textstellen bilden die Grenze klassischer maschineller Semantik: Absurdität, Ironie und semantische Brechung können von KI formal erkannt, aber nur eingeschränkt funktional übertragen werden.

Tab. 4: Übersetzungsbeispiel Märchenwelt

Originaltext	ChatGPT-Übersetzung	Übersetzung von Studierenden
<i>Fliegende Teppiche surrten durch die Luft, störrische Naturen, aber mit farbenfrohen Mustern</i>	У повітрі кружляли вперті летючі килими з барвистими візерунками	У повітрі кружляли летючі килими – уперті натури – з барвистими візерунками
<i>In den Schluchzschluchten hatte er mit einem Werwolfungen gerungen</i>	У ревучих ущелинах він боровся з юним вовкулаком	Боровся з перевертнем в ущелині, де ехом відлунювало ридання

Der erste Satz kombiniert Bewegung mit lautmalerischer Komponente (*surrten*), Fantastik (*fliegende Teppiche*) und ironisch-personifizierende Eigenschaften (*störrische Naturen*). Diese Mehrschichtigkeit ist für die Übersetzung zentral.

Die ChatGPT-Version integriert die Personifizierung durch das vorangestellte Adjektiv *вперті* direkt in die Nominalphrase. Die Struktur ist glattgezogen und stilistisch homogen, da das Modell meist nach Mustern optimaler Ukrainisch-Syntax sucht. Das Lexem *кружляли* beschreibt kreisende Bewegung, bildet jedoch nicht vollständig die lautmalerische Komponente von *surrten* ab. Die studentische Version versucht, die dreifache Segmentierung des deutschen Satzes durch Gedankenstriche nachzuahmen; *уперті натури*

überträgt *störrische Naturen* wörtlich und verstärkt die Personifizierung sichtbar. Die studentische Übersetzung ist stilistisch ambitionierter und bemüht sich, die rhetorische Form des deutschen Satzes sichtbar zu erhalten.

Im zweiten Satz überträgt die Phrase *у ревучих ущелинах* („in donnern- den/brüllenden Schluchten“) den akustischen Charakter der deutschen Wort- neuschöpfung, jedoch verschiebt sich die Emotionalität von *schluchzen* zu einem eher naturgewaltigen, lauten Geräusch (Rauschen, Brüllen). Die Zu- rückführung der Lautmalerei auf *ревучи* erzeugt ein dramatischeres, weniger wehmütiges Bild.

Die studentische Version entschlüsselt die metaphorische Struktur des Originals, indem sie *Schluchzschluchten* analytisch rekonstruiert: *в ущелині, де ехом відлунювало ридання* („in der Schlucht, wo Weinen als Echo wi- derhallte“). Dieser Ansatz fängt die im deutschen Kompositum enthaltene Ambivalenz aus *Raum* + *Schluchzen* deutlicher und differenzierter ein als in der ChatGPT-Fassung. Die studentische Version zeichnet sich durch eine poetischere und atmosphärisch reichere Rekonstruktion des Ausgangsbildes und die intensive Darstellung der akustisch-emotionalen Szenerie aus, die das deutsche Wortspiel *Schluchzschluchten* besonders wirkungsvoll transformiert. Damit zeigt die studentische Übersetzung eine höhere kreative Integration der Vorlage, während die maschinell orientierte Fassung stärker auf formale Treue und lexikalische Präzision setzt.

6 Fazit

Als Fazit der Studie lässt sich konstatieren, dass die studentische Teilnahme an den Übersetzungsprojekten die Entwicklung von Teamarbeits- und Redak- tionsfähigkeiten, die Erweiterung der sprach- und landeskundlichen Kom- petenz und die Steigerung der Motivation durch reale kreative Tätigkeiten förderte.

Die Arbeit mit ChatGPT regt zum Vergleich und zur Analyse von Überset- zungen an, ohne jedoch die kreative Verantwortung zu übernehmen. ChatGPT liefert Vorschläge, Formulierungen oder alternative Ansätze, die als Ausgangs-

punkt für Übersetzungen dienen können, aber die endgültige Interpretation, Bewertung und Umsetzung liegt stets bei den Studierenden. So bleibt der Lern- und Gestaltungsprozess menschlich geprägt, während KI unterstützend wirkt.

Die Unterrichtseinheit macht deutlich, dass KI-Tools im Bildungs- und Übersetzungskontext gleichermaßen Chancen und Herausforderungen mit sich bringen. Maschinelle Übersetzungen bieten schnelle, zugängliche und oft erstaunlich präzise Ergebnisse, doch werden ihre Grenzen – vor allem im Hinblick auf Stil, Pragmatik und kulturelle Nuancen – gerade im Vergleich mit menschlichen Lösungen sichtbar. Die gemeinsame Analyse verschiedener Systeme sowie die produktive Arbeit mit KI im literarischen Übersetzungsprozess zeigen, dass Künstliche Intelligenz als Impulsgeber fungieren kann, jedoch keinesfalls als Ersatz für menschliche Urteilskraft.

Besonders wertvoll erwies sich die Verbindung von maschineller Generierung, studentischer Reflexion und autorenbezogener Kontextualisierung. Sie verdeutlicht, wie wichtig es ist, KI nicht als automatischen Lösungsapparat, sondern als Werkzeug im Sinne einer didaktisch eingebetteten Übersetzungsstrategie zu nutzen. Das Projekt unterstreicht damit, dass die Zukunft der Übersetzung nicht im Entweder-oder von Mensch und Maschine liegt, sondern in einem bewussten, kompetenten und verantwortungsvollen Miteinander – getragen von kritischer Analyse, kreativem Denken und fundierter sprachlicher Expertise.

Literatur

- BDÜ (2023): „ChatGPT & Co.: Bedroht ‚Künstliche Intelligenz‘ die Berufe der Übersetzer und Dolmetscher?“ *Medieninformation des BDÜ vom 8.6.2023*. <https://bdue.de/fuer-presse-medien/presseinformationen/pm-detail/chatgpt-co-bedroht-kuenstliche-intelligenz-die-berufe-der-uebersetzer-und-dolmetscher> (5.12.2025).
- DOLMACI, MUSTAFA (2025): „Linguistic and Cultural Fidelity in AI-Driven Translation: A Comparative Analysis of DeepL and Human Translations in English–Turkish Literary Short Stories.“ *Uluslararası Dil Ve Çeviri Çalışmaları Dergisi* 5.1: 92–109. DOI: 10.63673/Lotus.1662868.

- ELKINS, KATHERINE (2024): „In Search of a Translator: Using AI to Evaluate what’s Lost in Translation.“ *Front. Comput. Sci.*, 13 August 2024. *Sec. Human-Media Interaction* 6. DOI: 10.3389/fcomp.2024.1444021.
- KARABAYEVA, IRINA / KALIZHANOVA, ANNA (2024): „Evaluating Machine Translation of Literature through Rhetorical Analysis.“ *Journal of Translation and Language Studies* 5.1: 1–9. DOI: 10.48185/jtls.v5i1.962.
- LANGE, ANJA / WOLF, DAGNY (2023): „Eine Fortbildung zu ChatGPT – Vorstellung des Unterrichtskonzepts und einer in Kirgistan durchgeführten Querschnittsstudie.“ *Germanistik in der Ukraine* 18: 24–32. DOI: 10.32589/2408-9885.2023.18.304860.
- MANAPBAYEVA, ZHANNURA / ZAURBEKOVA, GULKHADISHA / AYAZBEKOVA, KARLYGASH / KAZEZOVA, AIGUL / PIRMANOVA, KUNSULU (2024): „AI in Literary Translation: ChatGPT-4 vs. Professional Human Translation of Abai’s Poem ‚Spring.‘“ *Procedia Computer Science* 251: 526–531. DOI: 10.1016/j.procs.2024.11.143.
- NORD, CHRISTIANE (2011): *Funktionsgerechtigkeit und Loyalität: Die Übersetzung literarischer und religiöser Texte aus funktionaler Sicht*. TRANSÜD. Arbeiten zur Theorie und Praxis des Übersetzens und Dolmetschens. Berlin: Frank & Timme.
- UNESCO (2023): *Guidance for Generative AI in Education and Research*. Paris: UNESCO. <https://www.unesco.org/en/articles/guidance-generative-ai-education-and-research> (15.12.2025).

Loosening the Tightrope

Mindfulness, Phenomenology of Language and Embodied Listening in Interpreter Training

Abstract

Interpreting demands more than advanced cognitive skills: interpreters must follow and reformulate speech in real time amid fast delivery, variable accents and disfluent, non-ideal input. Yet training often relies on slow, standardized speech that encourages a narrow focus on semantic decoding and leaves students underprepared for the turbulence of real-life language use. Drawing on phenomenological accounts of language and mindfulness-based body–mind training, this chapter proposes an embodied pedagogy of listening centered on the preverbal stages of speech perception. Through short mindfulness practices and an adaptation of the Jean Georges Ernst Method, students learn to attune to rhythm, timbre and acoustic texture before moving to lexical and propositional meaning. The aim is to loosen the tightrope of idealized input and help future interpreters ride the waves of spoken variation, reclaiming human listening as a core professional competence in an age of AI-mediated language processing.

Keywords embodied listening, interpreter training, mindfulness, phenomenology of language, speech perception

1 Introduction: From cognitive overload to embodied presence

Interpreting is routinely described as a cognitively demanding activity. Descriptive and experimental studies portray conference, liaison and public service interpreting as high-load, high-pressure forms of language mediation that mobilize concurrent processes of listening and analysis, memory, speech production and self-monitoring (Gile 1995 [2009]; Pöchhacker 2004; Setton/Dawrant 2016). Training programs are therefore expected to expose students to tasks that approximate these conditions and prepare them for the realities of the booth or the field.

Classroom practice often reflects this orientation through a predominantly product-related understanding of quality. Teaching sequences are frequently organized around speed, accuracy and problem-solving: students are trained to identify key information, segment discourse, prioritize content and produce target-language reformulations as quickly and smoothly as possible. In this context, quality tends to be operationalized in terms of completeness and correctness, and feedback focuses on omissions, misinterpretations and terminological choices (Kalina 2002; Pöchhacker 2016). This sits uneasily with research on conference interpreting quality, which highlights its multi-dimensional, norm-governed nature and shows that quality is shaped by the expectations, roles and norms of the participants involved rather than by a single universal standard (Collados Aís et al. 2007; Pradas Macías/Zwischenberger 2022). In everyday teaching, however, this more complex picture is often reduced to what reaches the target language and how fast.

From the teacher's perspective, a different layer becomes visible. Many students appear to 'live in their heads': they mobilize intense mental effort yet show signs of anxiety, rigidity and self-doubt. When confronted with more authentic material — spontaneous speech, faster delivery, native and non-native accents — they often beg for slower, cleaner and 'less real' input, especially when working from their B language or any language other than their L1. In practice, they tend to chase meaning too fast, attempting to decode and reformulate content before they have fully stabilized the sounds they are hearing. Meanwhile, the interpreter's nervous system and lived bodily

experience are clearly affected, as studies on stress and fatigue have shown (Kurz 2002; Moser-Mercer 2005), but these dimensions rarely become explicit objects of pedagogy.

Listening in interpreter training is typically conceptualized as a decoding operation: the task is to “understand the message” and “transfer the meaning”. The interpreter’s body, although undeniably present, tends to be treated as a neutral container or as a technical concern limited to posture, breathing and vocal hygiene (Timarová/Salaets 2011). This neglect mirrors a broader Western tendency to privilege thinking over sensing: we are trained to think the language — its rules, categories and meanings — more than to feel its textures, rhythms and resonances. For interpreters, whose work unfolds in and through the living medium of speech, this asymmetry is striking. Ignoring the sounding, sensory dimension of language is not only a missed opportunity for well-being; it also constrains the quality and sustainability of interpreting performance in the long term.

In response to this gap, the present work proposes an embodied, mindfulness-informed approach to listening and interpreting. It brings together three strands: research on interpreting and attention (Gile 1995 [2009]; Pöchhacker 2004 [2016]; Kalina 2002); emerging work on mindfulness and interpreter training; and a specific listening protocol, the Jean Georges Ernst Method (JGEM), originally devised for foreign-language learning and here adapted to interpreter education. The contribution is primarily process-oriented, since it focuses on the phenomenology of listening and the training of embodied attention, but it is explicitly motivated by findings from product-oriented quality research on fluency, prosody and user perception. It seeks to show how refining interpreters’ preverbal listening processes can contribute to the very dimensions of quality that users actually hear and evaluate.

2 Interpreting, listening effort and mindfulness

Cognitive approaches have offered influential models for understanding the demands placed on interpreters. Gile's Effort Models conceptualize interpreting as the concurrent execution of distinct but interdependent efforts — listening and analysis, production, memory and coordination — within a limited pool of processing capacity (Gile 1995/2009). His tightrope hypothesis suggests that interpreters typically work near their maximum capacity, so that additional load (e.g. higher speech rate, unfamiliar accents, dense terminology) may tip the balance and lead to errors, omissions or breakdowns. Within this framework, the Listening and Analysis Effort is crucial: it encompasses “all comprehension-oriented operations, from the subconscious analysis of the sound waves carrying the source-language speech which reach the interpreter's ear through the identification of words to the final decisions about the ‘meaning’ of the utterance” (Gile 2009: 160). Drawing on psycholinguistic work, Gile recalls Guibert's (1979) finding that even trained phoneticians correctly identify no more than about 75% of phonetic segments on a sonogram (Gile 2009: 161), underscoring that speech recognition as it occurs in interpreting has important non-automatic components.

Since Gile's formulation, automatic speech recognition has improved dramatically, yet recent evaluations still show that state-of-the-art systems perform best with short, well-recorded speech and deteriorate with spontaneous conversational input, overlapping talk, strong accents or background noise (O'Connor Russell et al. 2024; Tobin et al. 2024). “AI interpreting” solutions, based on cascades of automatic speech recognition (ASR), machine translation and speech synthesis, continue to struggle with latency, omissions and interactional complexity (Nimdzi Insights 2022). In my own practice as an interpreter trainer, experimenting with Whisper-based tools, on-device ASR and cloud services such as Microsoft Azure Speech and Google Cloud Speech-to-Text, I have observed that these systems can handle high speech rates and accented speech remarkably well in not so favorable conditions, and they rarely break down in more spontaneous, noisy or emotionally charged situations. In

other words, even systems explicitly optimized for speech recognition do not treat the mapping from sound to words as trivial.

Work on interpreting competence has broadened this cognitive picture. (2004 [2016]) and Setton/Dawrant (2016) describe interpreting competence as a multidimensional construct combining linguistic, cognitive, technical and interpersonal components, while Kalina (2002) emphasizes quality as the outcome of strategic processing and prerequisites such as preparation, anticipation and mental readiness. Experimental work on quality perception in simultaneous interpreting, notably the Granada studies led by Collados Aís and colleagues, shows that manipulating single parameters such as accent, voice, intonation, style or grammatical correctness can significantly affect users' assessment of overall quality (Collados Aís et al. 2007). Users do not evaluate meaning in isolation: they respond to the way meaning is sounded. The approach proposed here takes these insights seriously and suggests that training at the level of preverbal, embodied listening can support both the cognitive demands described in processing models and the product-related dimensions of quality identified in user-based research.

In contemporary psychology and education, mindfulness is typically defined as the awareness that arises from paying attention, on purpose, in the present moment and non-judgmentally (Kabat-Zinn 1994). Mindfulness-based interventions such as Mindfulness-Based Stress Reduction and Mindfulness-Based Cognitive Therapy operationalize this principle through structured programs that combine formal meditation with informal exercises integrated into everyday life (Kabat-Zinn 1990; Segal et al. 2013). Rather than targeting cognition alone, these practices deliberately engage the whole organism — posture, breath, sensations, thoughts and emotions — in a systematic retraining of attention and awareness. Research on interoceptive awareness suggests that such practice can enhance people's ability to notice and make sense of subtle bodily cues, including tension, agitation and fatigue (Mehling et al. 2011, 2012). Many programs also include sound as an explicit object of practice, inviting participants to attend to the texture and transience of environmental sounds or spoken language (Kabat-Zinn 1994). For interpreters, whose work requires continuous management of attention and emotion in real time, these findings suggest that mindfulness can be more than a generic wellness tool:

it can become a deliberate component of professional training, aimed at supporting long-term quality and sustainability of practice.

Work explicitly addressing mindfulness in interpreter training is still at an early but promising stage. A small group of pioneering studies and student projects has explored mindfulness and related approaches in interpreting pedagogy, mainly by examining their effects on stress, attention and exam performance (Jiménez Ivars/Pinazo Calatayud 2013; Johnson 2016, 2022). Other contributions draw on acceptance-based or third-wave cognitive-behavioral frameworks to address performance anxiety, emotional strain and vicarious trauma — particularly in public service and dialogue interpreting — thus bringing questions of mental health and self-care into the heart of interpreter education (Hodáková/Ukušová 2022; Valverde-Vargas 2024), in line with broader developments in process-based cognitive behavioral therapy (CBT) (Hayes/Hofmann 2017). Together, these initiatives have been crucial in legitimizing mindfulness and acceptance-based tools as part of the methodological repertoire available to interpreter trainers. What is still missing, however, is a systematic, phenomenologically grounded protocol that focuses on listening itself, and more specifically on the preverbal stages of speech perception. The approach developed in this paper builds directly on these earlier achievements but shifts the focus towards using mindfulness not only for anxiety reduction or general cognitive enhancement, but also as a means of refining human oral speech-recognition skills at the preverbal level — consciously reclaiming, for human learners, forms of exposure-based training that have already been modelled computationally in neural systems.

3 The sounding body as pedagogical site

Phenomenological accounts of perception start from a simple but far-reaching claim: we do not first have abstract concepts and then apply them to a neutral stream of sense data; rather, we inhabit the world through a sensing, oriented body that is already engaged with its surroundings before any explicit interpretation takes place (Merleau-Ponty 1971). Perception is not a passive

reception of stimuli, but an active, bodily “taking up” of the world. This pre-reflective field of experience is structured, not chaotic: it already distinguishes figure from ground, intonation peaks from background hum, urgency from calm, even before words are clearly recognized.

Language, in this view, is not a disembodied code added on top of perception but one of the ways in which this embodied contact with the world becomes articulated. Phenomenological approaches to language emphasize that speech is a carnal, temporal phenomenon: words appear as sound-events, unfolding in time with a particular grain, rhythm and contour (Inkpin 2016; Hatab 2017). Children do not enter language by memorizing abstract rules; they plunge into a sea of voices, imitate sounds, play with syllables, taste words in the mouth. Meaning is inseparable from this sonic, gestural and affective texture. Eco-phenomenological reflections such as Abram’s meditation on the “magic of the senses” push this point further: the voice is not just a vehicle for information, but a vibratory exchange between bodies in a shared acoustic space (Abram 2000).

For interpreting, this has direct consequences. In the booth or in dialogue settings, interpreters are constantly immersed in streams of speech that arrive before they are fully intelligible. Accents, disfluencies, hesitations and overlapping turns all shape the acoustic field. At the same time, the interpreter’s own body is continuously responding: breathing changes, muscles tighten or release, the chest or throat resonates with certain voices, fatigue slowly accumulates. Usually, these responses remain in the background, noticed only when they become uncomfortable or disruptive. From a phenomenological perspective, however, this sounding body is not an incidental side effect of interpreting; it is part of the very medium through which listening and understanding take place (Merleau-Ponty 1971; Hatab 2019).

If we take this seriously, the listening effort described in cognitive models can be re-described, at least in part, as the work of stabilizing an acoustic world in and through the body: segmenting the flow of sound into rhythmic units, letting certain contours stand out, allowing familiar phonetic and prosodic patterns to emerge against noise. Before there are recognized words, there is already an embodied feel for what kind of utterance is coming — question, insistence, hesitation, joke — and for how manageable or threatening the

situation is. The proposal developed in the following sections starts precisely at this preverbal level. It treats the interpreter's sounding body as a legitimate object of training and suggests that, by systematically cultivating awareness of rhythm, resonance and acoustic texture, it is possible to support the more classical goals of interpreting pedagogy: reliable word recognition, stable attention and, ultimately, a spoken product that users hear as fluent, clear and trustworthy.

4 The Jean Georges Ernst Method: training preverbal listening

The Jean Georges Ernst Method (JGEM) was originally devised for foreign-language learning, not for interpreting. Its core intuition is simple but powerful: before understanding, we must learn to hear; before words, there are rhythms, contours and textures of sound. Instead of starting from written text or explicit grammar, the method starts from the acoustic flow itself and slows it down just enough to make its internal structure tangible.

In biographical terms, the method reached me through a German teacher in Geneva whose classes radically transformed the way I learned languages. Rather than explaining rules or handing out texts, he invited us to work with short audio recordings, sometimes only a few seconds long, and to reconstruct them step by step from their rhythmic skeleton to the eventual words. That experience later became the seed for the proposal developed here: adapting the perceptive phase of the JGEM to interpreter education and linking it explicitly with mindfulness and phenomenology.

In its basic version, the perceptive phase of the JGEM can be summarized in three steps:

1. Rhythmic segmentation (ta-lines): A short fragment of spoken language is played once, then again. On the second listening, learners pause the audio at points they intuitively perceive as sequences and reproduce only the rhythm and intonation with a neutral

syllable such as *ta*, without trying to imitate individual sounds. Each sequence is then recorded as a numbered line of *ta* units (or short underscores), creating a simple visual score of the utterance's temporal structure.

2. Inner echo and voice release: Once the *ta*-lines are in place, the same fragment is played again. Learners are explicitly instructed not to repeat immediately. First, they keep the tongue still and allow the sequence to resonate silently inside the body: feeling how the cadence, intonation and pauses vibrate in the chest, throat, jaw or face, noticing where the breath would naturally fall and registering any tension or ease that appears. Only after this inner echo has been fully allowed do they move on to the repetition phase, giving voice to what has just resonated internally. The repetition should come directly out of the echo — almost as a release of what the body has just held — rather than as a fresh attempt to get it right. Skipping the echo is not permitted in this step, because it is precisely in that brief, silent interval that learners assimilate and embody the sound before sending it back into the world.
3. From sound to grapheme: Immediately after repeating, learners write down what they think they have just said, placing graphemes above the *ta*-lines without worrying about correctness or standard phonology. Approximations, invented spellings and partially recognized clusters are allowed—and even desirable. The aim is to give a concrete written trace to auditory impressions that would otherwise remain fleeting.

For interpreting, two features of this protocol are particularly relevant. First, it makes visible processes that usually remain hidden: the segmentation of the acoustic flow into manageable units, the negotiation between rhythm and phonetic detail, the moment where noise becomes words. Second, it legitimizes a preverbal stage of listening practice in which students are not yet required to understand or translate, but simply to stay with the complexity of the speech signal. This runs counter to the dominant classroom reflex of going straight

to meaning and can be unsettling for students who are used to being assessed on immediate comprehension and reformulation.

In my own classes on French–Spanish interpreting in commercial and touristic settings, I have introduced this perceptive phase in a single two-hour session. Initial reactions have often been skeptical: students question the usefulness of exercises that delay understanding and do not seem to train interpreting in the conventional sense. At the same time, once they have gone through several cycles of *ta-lines*, inner echo and voice release, and grapheme reconstruction, most report being able to imitate accents and speech rates they would normally find threatening, and to recognize words and expressions they had previously perceived as an undifferentiated blur. Even after such a brief intervention, many describe a reduction in stress when exposed to unfamiliar accents, together with a clearer sense of how their own body responds to fast or non-standard speech. These are, of course, qualitative impressions rather than experimental results, but they suggest that the JGEM perceptive phase can open a pedagogical space that is rarely addressed in interpreter training: the slow, patient cultivation of listening before understanding.

5 Integrating mindfulness and JGEM in interpreter training

The proposal developed in this paper is to integrate short mindfulness practices with the perceptive phase of the Jean Georges Ernst Method within interpreter training. The aim is not to turn interpreting classes into meditation courses, but to use a small number of targeted practices to stabilize attention, refine bodily awareness and support the kind of sustained, non-reactive listening that the JGEM exercises require.

A basic teaching sequence might have three phases:

1. Settling the body and attention (5–10 minutes): The session begins with a brief mindfulness practice, either seated or standing. Students are invited to adopt a balanced posture, notice points of contact with

the chair or floor, and bring gentle attention to the breath. The focus is not on achieving a particular state, but on recognizing sensations of tension or restlessness and allowing them to be present without immediate reaction. A short sound-based practice can follow, in which students simply listen to the ambient sounds of the room or corridor, noticing changes in volume, distance and timbre.

2. Perceptive work with speech (20–30 minutes): The JGEM perceptive phase is then applied to short fragments of authentic spoken material in the working languages (e.g. segments of speeches, interviews or dialogues). Students move through the three steps described above, from ta-lines to emerging words, while teachers periodically bring their attention back to bodily cues: “What do you notice in your shoulders as the speech accelerates?”, “What happens in your breathing when you cannot catch a word?”, “Can you stay with the rhythm even when you feel impatient?”. In this way, mindfulness becomes a lens for observing not only the speech signal but also the interpreter’s internal reactions to it.
3. Linking to interpreting tasks (20–30 minutes): Finally, the same or similar fragments are used for short interpreting exercises (e.g. shadowing, summarized oral translation, short consecutive or bilateral dialogue). Students are encouraged to notice whether their relationship to the input has shifted: Do they feel more comfortable with accented speech? Can they tolerate ambiguity for a little longer before committing to a reformulation? Is it easier to remain present when they miss a word? Brief reflective discussions or written logs can help articulate these observations and connect them with students’ developing sense of professional quality.

The integration of mindfulness and JGEM thus operates on two levels. At a micro-level, mindfulness practices support the perceptive work by helping students stay with the difficulty of fast, accented or messy speech without collapsing into anxiety or self-criticism. At a macro-level, the combined protocol functions as a process-oriented quality intervention: instead of training only the visible product of interpreting, it seeks to strengthen the

underlying listening processes that feed into fluency, prosody and credibility—precisely the parameters that users tend to notice and evaluate.

The format is flexible. Short JGEM-mindfulness modules can be inserted periodically into existing courses, for example as a recurring listening laboratory strand or concentrated into intensive workshops at particular stages of training. They can also be adapted to different language combinations and modes (simultaneous, consecutive, bilateral), provided that teachers select material whose length and complexity match the students' current level. What matters is not the specific texts chosen, but the pedagogical stance: listening treated as an embodied, trainable practice in its own right, rather than a transparent prelude to 'real interpreting'.

6 Discussion and conclusion: towards an embodied pedagogy of listening

The argument developed in this contribution can be summarized in three moves. First, drawing on cognitive models and quality research, it has highlighted that interpreting is already demanding at the level of listening and speech recognition, before higher-level discourse processing comes into play (Collados Aís et al. 2007; Gile 2009). Users respond strongly to how speech sounds, not only to what it means, and even advanced ASR systems continue to struggle with the full complexity of natural, accented and noisy speech (Nimdzi Insights 2022; O'Connor Russell et al. 2024; Tobin et al. 2024). Second, building on phenomenology and mindfulness research, it has been argued that listening is not a purely cognitive act but an embodied encounter in which the interpreter's sounding body plays an active role (Merleau-Ponty 1971; Kabat-Zinn 1990; Abram 2000; Hatab 2017). Third, it has proposed the perceptive phase of the Jean Georges Ernst Method, integrated with short mindfulness practices, as a concrete way of training this preverbal, embodied layer of listening within interpreter education.

The empirical status of this proposal is deliberately modest. My own classroom experience at the University of Granada, particularly in French–

Spanish bilateral interpreting for commercial and touristic settings, suggests that even brief exposure to the JGEM perceptive phase can change how students relate to challenging input: they become more capable of imitating non-standard accents, more tolerant of temporary incomprehension and more aware of their bodily reactions to fast or dense speech. At the same time, resistance and impatience are common, especially when students feel that they are ‘not really interpreting’ because they are not yet producing target-language output. This points to a cultural challenge within interpreter training itself: the need to legitimize slow, non-productive forms of practice that work beneath the surface of performance.

Several limitations need to be acknowledged. The activities described here have so far been implemented in small groups, over very short periods and without control groups or standardized measures. No claims can be made about their efficacy beyond the qualitative impressions reported. Mindfulness practices themselves are not neutral: they require careful framing to avoid becoming a tool for merely tolerating excessive workloads or structural problems in working conditions (Purser 2019). Any attempt to integrate them into professional education should therefore be accompanied by critical discussion of their aims, limits and institutional context.

These caveats suggest directions for future research. Small-scale experimental studies could compare classes with and without JGEM-mindfulness modules, using mixed methods that combine performance measures, anxiety scales and qualitative reflection. Longitudinal designs might explore whether early training in preverbal listening has lasting effects on students’ tolerance of accented speech, perceived stress and quality of delivery. Collaborative work between interpreter trainers, phenomenologists of language and mindfulness researchers could further refine the conceptual framework and tailor practices to the specific constraints of interpreting classrooms.

Despite these limitations, the proposal points towards a broader pedagogical shift. If interpreting is to remain a sustainable and high-quality practice in a landscape increasingly shaped by AI, it may not be enough to optimize terminology databases, booth technology or external working conditions. Attention also needs to turn inward, towards the listening lives of interpreters themselves: how they inhabit their bodies while listening, how

they meet the acoustic complexity of real speech, and how they recover when their efforts fail. An embodied pedagogy of listening, informed by mindfulness and phenomenology and operationalized through methods such as the JGEM perceptive phase, invites us to train interpreters not only as fast problem-solvers, but as finely tuned listeners who can photograph language in motion because mind and sounding body are working together, rather than against each other. To borrow Gile's (2009) image, interpreters are often described as walking a tightrope at the limits of their capacity, clinging to balance while any change in speed, accent or register threatens a fall. The approach outlined here does not deny the height of that rope; it simply refuses to treat variation as the enemy. Its ambition is to loosen the cable, let it dip into the swell of living speech, and teach students to ride the movement instead of freezing before it. In a world where spoken language behaves less like a straight line and more like an ocean — turbulent, accented, improvised — quality may depend less on reproducing an ideal script than on cultivating the courage and sensitivity to surf those waves without losing ethical or semantic bearings. If this proposal has a single wager, it is that interpreters who learn to inhabit their sounding bodies in this way will not only survive the tightrope of professional practice; they may also help redraw it in more humane, sustainable and imaginatively responsive ways.

References

- ABRAM, DAVID (2000): *La magia de los sentidos*. Barcelona: Kairós.
- BARBEZAT, DANIEL P. / BUSH, MIRABAI (2013): *Contemplative Practices in Higher Education: Powerful Methods to Transform Teaching and Learning*. San Francisco: Jossey-Bass.
- COLLADOS AÍ, ÁNGELA / PRADAS MACÍAS, M. MACARENA / STÉVAUX, ELISABETH / GARCÍA BECERRA, OLALLA, EDS. (2007): *La evaluación de la calidad en interpretación simultánea: parámetros de incidencia*. Granada: Comares.
- GILE, DANIEL (1995 [2009]): *Basic Concepts and Models for Interpreter and Translator Training*. Amsterdam/Philadelphia: Benjamins.
- GUIBERT, JEAN (1979): *La parole. Compréhension et synthèse de la parole par les ordinateurs*. Paris: Presses Universitaires de France.

- HATAB, LAWRENCE J. (2017): *Proto-Phenomenology and the Nature of Language: Dwelling in Speech I*. London: Rowman & Littlefield International.
- HATAB, LAWRENCE J. (2019): *Proto-Phenomenology, Language Acquisition, Orality, and Literacy: Dwelling in Speech II*. London: Rowman & Littlefield International.
- HAYES, STEVEN C. / HOFMANN, STEFAN G. (2017): "The Third Wave of Cognitive Behavioural Therapy and the Rise of Process-Based Care." *World Psychiatry* 16.3: 245–252.
- HODÁKOVÁ, SOŇA / UKUŠOVÁ, JANA (2022): "Mental Health in Public Service Interpreting." *BRIDGE: Trends and Traditions in Translation and Interpreting Studies* 3.1: 58–70.
- INKPIN, ANDREW (2016): *Disclosing the World: On the Phenomenology of Language*. Cambridge, MA: MIT Press.
- JENNINGS, PATRICIA A. (2015): *Mindfulness for Teachers: Simple Skills for Peace and Productivity in the Classroom*. New York: Norton.
- JIMÉNEZ IVARS, AMPARO / PINAZO CALATAYUD, DANIEL (2013): "Mindfulness Training for Interpreting Students." *Lebende Sprachen* 58.2: 341–65. DOI: 10.1515/les-2013-0020.
- JOHNSON, JULIE E. (2016): *Effect of Mindfulness Training on Interpretation Exam Performance in Graduate Students in Interpreting*. Doctoral dissertation. San Francisco: University of San Francisco.
- JOHNSON, JULIE E. (2022): "Mindfulness Training for Conference Interpreters." *The Routledge Handbook of Conference Interpreting*. Eds. M. Albl-Mikasa / E. Tiselius. London / New York: Routledge. 564–580. DOI: 10.4324/9780429297878-48.
- KABAT-ZINN, JON (1990): *Full Catastrophe Living: Using the Wisdom of Your Body and Mind to Face Stress, Pain, and Illness*. New York: Delacorte.
- KABAT-ZINN, JON (1994): *Wherever You Go, There You Are: Mindfulness Meditation in Everyday Life*. New York: Hyperion.
- KALINA, SYLVIA (2002): "Quality in Interpreting and its Prerequisites: A Framework for a Comprehensive View." *Interpreting in the 21st Century: Challenges and Opportunities*. Eds. Giuliana Garzone / Maurizio Viezzi. Amsterdam/Philadelphia: Benjamins. 21–130.
- KURZ, INGRID (2002): "Physiological Stress Responses during Media and Conference Interpreting." *Benjamins Translation Library* 43. Eds. Maurizio Viezzi / Giuliana Garzone. 195–202. DOI: 10.1075/btl.43.19kur.

- MA, S. HELEN / TEASDALE, JOHN D. (2004): "Mindfulness-Based Cognitive Therapy for Depression: Replication and Exploration of Differential Relapse Prevention Effects." *Journal of Consulting and Clinical Psychology* 72.1: 31–40. DOI: 10.1037/0022-006X.72.1.31.
- MEHLING, WOLF E. / WRUBEL, JUDITH / DAUBENMIER, JENNIFER J. / PRICE, CYNTHIA J. / KERR, CATHERINE E. / SILOW, TRACY / GOPISETTY, VASANTHA / STEWART, ANITA L. (2011): "Body Awareness: A Phenomenological Inquiry into the Common Ground of Mind-Body Therapies." *Philosophy, Ethics, and Humanities in Medicine: PEHM* 6.1: 6. DOI: 10.1186/1747-5341-6-6.
- MEHLING, WOLF E. / PRICE, CYNTHIA / DAUBENMIER, JENNIFER J. / ACREE, MICHAEL / BARTMESS, ERIN / STEWART, ANITA L. (2012): "The Multidimensional Assessment of Interoceptive Awareness (MAIA)." *PLOS ONE* 7.11: e48230. DOI: 10.1371/journal.pone.0048230.
- MERLEAU-PONTY, MAURICE (1971): *Phénoménologie de la perception*. (Nouv. tirage). Paris: Gallimard.
- MOSER-MERCER, BARBARA (2005): "Remote Interpreting: Issues of Multi-Sensory Integration and Cognitive Load." *Meta (Montréal)* 50.2: 727–738. DOI: 10.7202/011014ar.
- NIMDZI INSIGHTS (2022): *Machine Interpreting (MI): Your Ultimate Guide on Industry Landscape and Solutions*. Seattle: Nimdzi Insights. <https://www.nimdzi.com/machine-interpreting-mt/> (29.11.2025).
- O'CONNOR RUSSELL, SARAH / GESSINGER, INGO / KRASON, AGNIESZKA / VIGLIOCO, GABRIELLA / HARTE, NAOMI (2024): "What Automatic Speech Recognition Can and Cannot Do for Conversational Speech Transcription." *Research Methods in Applied Linguistics* 3.3: 100163. DOI: 10.1016/j.rmal.2024.100163.
- PÖCHHACKER, FRANZ (2004): *Introducing Interpreting Studies*. London / New York: Routledge.
- PÖCHHACKER, FRANZ (2016): *Introducing Interpreting Studies*. 2nd ed. London / New York: Routledge.
- PRADAS MACÍAS, E. MACARENA / ZWISCHENBERGER, CORNELIA (2022): "Quality and Norms in Conference Interpreting." *The Routledge Handbook of Conference Interpreting*. Eds. Michaela Albl-Mikasa / Elisabet Tiseliuss. London / New York: Routledge. 243–257. DOI: 10.4324/9780429297878-24.

- PURSER, RONALD (2019): *McMindfulness: How Mindfulness Became the New Capitalist Spirituality*. London: Watkins Media.
- SEGAL, ZINDEL V. / WILLIAMS, J. MARK G. / TEASDALE, JOHN D. (2013): *Mindfulness-Based Cognitive Therapy for Depression*. 2nd ed. New York: Guilford.
- SETTON, ROBIN / DAWRANT, ANDREW (2016): *Conference Interpreting – A Trainer’s Guide*. Amsterdam/Philadelphia: Benjamins.
- TIMAROVÁ, ŠÁRKA / SALAETS, HEIDI (2011): “Learning Styles, Motivation and Cognitive Flexibility in Interpreter Training: Self-Selection and Aptitude.” *Aptitude for Interpreting*. Eds. Franz Pöchhacker / Mabel Liu. Amsterdam/Philadelphia: Benjamins. 33–54.
- TOBIN, JOSEPH / NELSON, PAIGE / MACDONALD, BRENT / HEYWOOD, RACHEL / CAVE, RACHEL / SEAVER, KATHERINE / DESJARDINS, AMY / JIANG, PO-PING / GREEN, JORDAN R. (2024): “Automatic Speech Recognition of Conversational Speech in Individuals with Disordered Speech.” *Journal of Speech, Language, and Hearing Research* 67.11: 4176–4185. DOI: 10.1044/2024_JSLHR-24-00045.
- VALVERDE-VARGAS, SOFÍA (2024): “Stress, Vicarious Trauma and Emotional Impact of Dialogue Interpreters in Healthcare Settings: An Exploratory Review of Coping Strategies and Self-Care.” *Abant Journal of Translation and Interpreting Studies* 2.2: 64–70.
- WILBER, KEN (2016): *Integral Meditation: Mindfulness as a Way to Grow Up, Wake Up, and Show Up in Your Life*. Boulder, CO: Shambhala.
- ZAJONC, ARTHUR (2013): “Contemplative Pedagogy: A Quiet Revolution in Higher Education.” *New Directions for Teaching and Learning* 134: 83–94. DOI: 10.1002/tl.20057.

New Trends in the Interpreter's Toolbox

AI's Impact on Pre-Interpreting Tasks

Abstract

Die rasante Entwicklung der Künstlichen Intelligenz verändert auch das Konferenzdolmetschen grundlegend. Diese Studie untersucht den Einsatz von KI-gestützten Tools in der Vorbereitungsphase und deren Auswirkungen auf Berufs- und Studierendolmetscher. Auf der Grundlage eines Mixed-Methods-Designs werden sowohl Umfragen zu aktuellen Praktiken und Einstellungen als auch angeleitete Übungen analysiert, in denen der Einsatz von KI vor und nach gezieltem Training dokumentiert wird. Die Studie kartiert relevante Tools und vergleicht KI-unterstützte Terminologiesuchen, Glossarerstellung und Textzusammenfassungen mit traditionellen Methoden. Die Ergebnisse zeigen deutliche Vorteile der Künstlichen Intelligenz in Bezug auf Geschwindigkeit und Zugang zu spezialisiertem Wissen, betonen jedoch zugleich die Notwendigkeit menschlicher Kontrolle zur Sicherung von Genauigkeit und Kontextangemessenheit. Die Teilnehmenden erkennen die Vorteile an, äußern aber Bedenken hinsichtlich Datenschutz, Qualitätsunterschieden und möglicher Überabhängigkeit. Die Studie plädiert für eine ausgewogene Integration: KI soll menschliche Expertise ergänzen, nicht ersetzen. Abschließend werden Empfehlungen für Ausbildungsprogramme formuliert, die KI-Kompetenz, ethisches Bewusstsein und kritisches Urteilsvermögen als zentrale Fähigkeiten verankern.

Keywords Artificial Intelligence, conference interpreting, digital transformation, interpreter training, terminology management

1 Artificial Intelligence and the transformation of interpreting practice

The interpreting profession is experiencing a significant paradigm shift as Artificial Intelligence (AI) becomes an integral part of its ecosystem. Historically defined as a cognitively demanding activity dependent on linguistic expertise, domain knowledge, and strategic pre-task preparation (Gile 2009; Kalina 2015), interpreting has long relied on meticulous preparatory routines aimed at mitigating the cognitive load of live performance. These routines — comprising documentation review, terminology extraction, and conceptual mapping — enable interpreters to anticipate discourse, manage processing capacity, and maintain accuracy under pressure. In recent years, however, advances in machine learning, natural language processing (NLP), and large language models have begun to reshape these foundational processes, marking a new stage in the evolution of the field.

Contemporary interpreters now have access to a wide range of intelligent tools such as ChatGPT, Microsoft Copilot, or DeepL, which support automatic terminology extraction, document summarisation, and predictive analysis of content (Liu/Cheung 2023; Chen/Kruger 2024). These systems extend far beyond translation workflows, positioning AI as an active co-pilot that enhances, rather than replaces, professional expertise. Yet, this transformation is not purely technical: it entails a redefinition of the interpreter's agency, requiring practitioners to integrate automation into cognitively complex tasks without eroding the analytical depth, ethical awareness, and situational adaptability that constitute the core of the profession.

1.1 From translation technology to hybrid interpreting workflows

The historical trajectory of AI in language services — from early rule-based systems to neural and generative architectures — illustrates a gradual expansion in both functionality and cognitive reach (Downie 2020). Neural Machine Translation and transformer-based models such as BERT and GPT introduced

semantic depth and contextual coherence into machine output (Mellinger/Hanson 2018). Although initially used in translation, these architectures have rapidly migrated into interpreting, where NLP algorithms now enable the extraction of domain-specific terminology, the identification of named entities, and the clustering of recurrent concepts across large textual corpora (Álvarez Pérez 2024).

Additionally, the emergence of hybrid workflows is materialised in platforms such as KUDO or Interprefy, which include a whole range of features that illustrate an increasingly symbiotic relationship between human and machine, one that redistributes cognitive demands and procedural control. As Fantinuoli (2023) observes, AI is not a replacement but an amplifier of human capacity, enabling interpreters to transition from reactive problem-solving to proactive knowledge management.

1.2 AI and pre-task preparation

Among all phases of interpreting, pre-task preparation has been most visibly transformed by the integration of AI. Empirical evidence from the studies discussed later in further detail indicates that professionals deploy AI predominantly for terminology retrieval, glossary elaboration, and document summarisation — activities that are both time-consuming and cognitively demanding (Fantinuoli/Montecchio 2022). The adoption of intelligent tools in these areas has reshaped the temporal and cognitive architecture of preparation, redefining how interpreters acquire and organise domain-specific knowledge.

Applications such as InterpretBank or Intragloss employ NLP and statistical modelling to extract salient terms automatically, allowing interpreters to refine and verify them according to context (Fantinuoli 2023). Similarly, summarisation algorithms identify key arguments and thematic structures across large corpora, enabling interpreters to prioritise relevant information efficiently (Zhang et al. 2023). Automatic speech recognition-based systems add a performative dimension by transcribing recorded materials, synchronising them with glossaries, and thus facilitating rehearsal under quasi-authentic

conditions (Mellinger/Hanson 2018). Collectively, these technologies shift preparation from a predominantly manual, linear process to a dynamic cycle of machine assistance and human validation.

The benefits reported by interpreters are consistent: greater efficiency, faster access to terminological resources, and improved organisation of preparatory materials. However, these advantages are tempered by concerns regarding reliability, depth, and data security. Practitioners emphasise the need for critical oversight of AI-generated outputs, noting that terminological precision and contextual adequacy still depend on expert intervention. Ethical considerations are also prominent; the use of confidential documents in commercial AI environments raises questions about privacy and professional responsibility (Defrancq/Corpas Pastor 2023). Most interpreters therefore favour a complementary model that merges automation with traditional methods, using AI for acceleration while preserving human judgement for verification and conceptual integration. This approach resonates with Gile's (2009) Effort Models, which underscore the balance between reducing cognitive load and maintaining comprehension depth.

In short, AI has inaugurated a new form of preparatory ecology — one that combines algorithmic scalability with human interpretation. The interpreter's role increasingly resembles that of an editor or curator who transforms machine output into cognitively usable knowledge. Efficiency and reflection are no longer opposing values but interdependent dimensions of professional competence.

2 Artificial Intelligence in interpreter training

As AI reshapes professional practice, interpreter education must evolve in tandem. The proliferation of intelligent tools in the classroom demands not only technical adoption but also critical theorisation of their pedagogical role. Training programmes can no longer treat AI as an ancillary resource; they must incorporate it as an integral component of professional formation, while

ensuring that automation serves the development of cognitive autonomy and ethical discernment (Defrancq/Corpas Pastor 2023).

2.1 Pedagogical integration and emerging models

AI-driven technologies — particularly computer-assisted interpreting (CAI), ASR, and corpus-analysis tools — are increasingly embedded in interpreter training. They facilitate the creation of personalised glossaries, automate feedback through transcription analysis, and simulate realistic interpreting conditions (Chen/Kruger 2024). The pedagogical possibilities offered by these tools are unprecedented, yet their incorporation often precedes the necessary reflection on learning design. The pandemic context accelerated this process by normalising remote instruction and cloud-based communication (Liu 2022), creating digital infrastructures that now persist as standard practice.

Against this backdrop, Antonova (2023) introduces the notion of the *augmented interpreter*: a professional capable of collaborating with AI while preserving critical judgment and communicative autonomy. This concept situates technological competence within a broader epistemic and ethical framework. The effective interpreter of the future will not merely operate digital tools but will understand their affordances, constraints, and implications for cognition and ethics. Consequently, interpreter education must expand its scope to cultivate critical digital literacy, encompassing prompt formulation, data management, algorithmic transparency, and bias awareness. As Bhattacharya (2025) notes, such competence must also extend to educators themselves: without digitally literate instructors, the integration of AI into curricula risks devolving into uncritical mimicry of technological trends rather than informed pedagogical innovation.

2.2 Pedagogical benefits and cognitive gains

Empirical research increasingly substantiates the pedagogical value of AI when integrated thoughtfully. Corpus-based preparation workflows have been

shown to improve terminological recall and reduce omissions in simultaneous interpreting tasks (Xu 2018). By automating data collection, these tools free cognitive resources that can be allocated to higher-order processes such as discourse analysis and reformulation (Yang/Mu 2024). Moreover, ASR-generated transcripts allow for objective, time-stamped analyses of performance, transforming transient oral outputs into tangible learning artefacts (Chen/Kruger 2024). CAI platforms like InterpretBank further enhance realism by allowing students to rehearse in simulated environments with adaptive parameters such as speech rate, accent diversity, and topic complexity (Sun 2025).

2.3 Risks and constraints of AI integration

Despite its pedagogical promise, AI integration also entails cognitive, ethical, and epistemological risks that require careful management. When inadequately scaffolded, the multimodal complexity of AI environments can produce cognitive overload. Students confronted with simultaneous streams of information — speech, captions, terminology prompts — may experience fragmented attention and reduced comprehension (Guo et al. 2023). Similarly, the convenience of automation can foster dependency: overreliance on AI-generated glossaries or summaries risks weakening the analytical and reformulative capacities that are essential for professional interpreting (Liu 2022; Budiharjo 2024).

Ethical concerns are equally significant. The uploading of authentic or client-related materials to commercial platforms raises questions of confidentiality and data protection, while algorithmic bias can distort linguistic representation, particularly for under-resourced languages or non-standard accents (Zhang 2017). These issues highlight the necessity of embedding AI within a critical pedagogical framework that foregrounds responsibility and reflective decision-making. As Fan (2024) argues, genuine digital literacy involves not only technical competence but also the ability to discern when automation enhances performance and when it compromises interpretive integrity. The pedagogical objective should therefore be to train interpreters

capable of managing the tension between technological efficiency and human judgment, understanding that mastery of AI entails both its utilisation and its restraint.

2.4 Curriculum development and pedagogical outlook

The curricular implications of these transformations are far-reaching. AI literacy should be introduced early and reinforced across all stages of interpreter education, ensuring progressive acquisition of both operational and critical competences. Students must learn to design coherent workflows that combine manual and automated preparation, developing strategies for verification, revision, and ethical compliance. Equally important is the cultivation of an ethical mindset that regards data privacy, bias mitigation, and environmental sustainability as intrinsic components of professional responsibility.

Educators, for their part, must undergo parallel training to model reflective and responsible AI use. Teaching should not only demonstrate best practices but also expose limitations and failures, helping students to diagnose, correct, and learn from them. In this way, technological literacy becomes inseparable from the development of professional maturity. Ultimately, AI integration should lead not to dependency but to *augmentation* — the reinforcement of human agency through the intelligent use of digital tools.

3 Methodology

To address the dual focus of this research — professional interpreters and interpreting students — a mixed-methods design was adopted to capture both quantitative and qualitative dimensions of how Artificial Intelligence (AI) tools are used, perceived, and evaluated at different stages of expertise. Quantitative data were collected through structured questionnaires, while open-ended responses and guided classroom exercises provided qualitative depth. This

methodological integration enabled a comprehensive analysis that combined cross-sectional observations of professional practice with longitudinal insights into students' evolving engagement with AI.

3.1 Research design

A convergent mixed-methods approach (Creswell/Plano Clark 2018) was employed to triangulate numerical data and descriptive evidence. The study was organised into two complementary components.

- **Study 1** comprised a quantitative survey of professional interpreters designed to assess current patterns of AI use, the perceived reliability of AI-generated resources, and expectations concerning its future role in interpreting.
- **Study 2** involved a two-phase pedagogical intervention with undergraduate interpreting students, examining how guided exposure to Microsoft Copilot influenced their confidence, trust, and performance in pre-task preparation.

This combined design yielded a multidimensional understanding of AI integration: Study 1 captured how practising interpreters incorporate AI into their established workflows, while Study 2 examined the formative impact of structured training on student interpreters' developing technological and cognitive competences.

3.2 Participants

Professional interpreters (Study 1)

Forty-three practising interpreters participated in the survey (30 women, 13 men). The sample represented a wide age and experience spectrum: sixteen participants were aged 41–60, fifteen were between 26 and 40, seven were 20–25, and five were over 60. Professional experience ranged from beginners

(0–5 years; $n = 11$) to experts with more than two decades in the field ($n = 15$). This heterogeneity provided a balanced view of how interpreters at different career stages perceive and utilise AI in their professional routines.

Interpreting students (Study 2)

Thirty-seven undergraduates enrolled in a third-year Consecutive Interpreting BII (English) course within a university Translation and Interpreting programme participated in at least one phase of the study. Phase 1 included 33 students (mean age = 20.52) and phase 2 included 29 (mean age = 20.79), with ages clustering between 20 and 22. Some participants completed both phases, allowing for partially overlapping subgroups and enabling both within-subject and between-subject comparisons.

3.3 Instruments

3.3.1 Professional interpreters

The instrument for Study 1 combined Likert-scale items with open-ended questions to explore interpreters' use of AI in pre-task preparation. It included sections on general attitudes toward AI, frequency of use, and perceived complementarity with traditional preparation methods. Additional items examined specific applications — terminology retrieval, glossary compilation, document analysis, and summarization — as well as perceptions of efficiency, reliability, and quality compared to manually created resources. The final section addressed usability, data privacy, training needs, and ethical concerns, while open-ended items encouraged participants to elaborate on advantages, limitations, and professional expectations regarding AI integration.

3.3.2 Interpreting students

In Study 2, students completed two questionnaires — one administered before and one after a guided exercise with Microsoft Copilot. The pre-task questionnaire gathered baseline data on their familiarity with AI tools and chatbots, previous experiences, and expectations concerning potential benefits, ethical issues, and professional relevance. The post-task questionnaire

explored prompting strategies, evaluations of Copilot outputs, and perceptions of its utility in preparation. It also incorporated statements adapted from the Technology Acceptance Model to assess perceived ease of use, usefulness, attitudes toward AI, and intentions to continue employing such tools. Finally, students compared preparation with and without Copilot, reflecting on efficiency, learning depth, and any challenges encountered during the activity.

3.3.3 Procedure

For Study 1, the questionnaire was distributed online and completed anonymously. Participation was voluntary, and no identifying information was collected. Respondents answered individually and asynchronously, ensuring confidentiality and flexibility.

For Study 2, data collection was conducted in two consecutive classroom sessions designed to replicate authentic pre-interpreting conditions. In the first phase, students completed a preparatory task manually, without AI assistance, and filled in the pre-questionnaire. Several weeks later, they repeated a comparable task using Microsoft Copilot and then completed the post-questionnaire reflecting on their experience. All sessions adhered to institutional ethical standards and data protection protocols. Microsoft Copilot was used as the officially endorsed AI platform, ensuring a secure, standardised environment consistent with the educational context.

4 Data analysis

The analytical strategy combined descriptive statistics — means, medians, and frequency distributions — with an inductive thematic analysis of open-ended responses (Braun/Clarke 2006). In Study 1, Likert-type items captured general attitudes, frequency of use, and evaluations of AI tools among professional interpreters. In Study 2, within-group comparisons of pre- and post-task questionnaires assessed changes in perceived usefulness, confidence, and trust following guided exposure to Microsoft Copilot. Qualitative evidence refined these results by documenting perceived efficiency, trust, usability, ethical

concerns, limitations, and, in the student cohort, the evolution of prompting strategies and difficulties encountered during tool interaction. Triangulation of quantitative and qualitative strands provides an integrated account of how novice and experienced interpreters perceive and experience AI in pre-task preparation.

5 Professional interpreters: results and discussion

Survey data show limited commitment to paid services. A substantial majority (74.4%) reported no subscription to AI applications. One in five (20.9%) used ChatGPT Plus; one participant (2.3%) reported Gemini; none paid for Microsoft Copilot; 14% cited other platforms, and a single respondent reported multiple subscriptions. This pattern suggests experimentation constrained by cost–benefit considerations and scepticism about the added value of premium features, consistent with a phase of strategic rather than full adoption.

Perceived usefulness is, nevertheless, notable. Almost half of the sample (48.8%) rated AI as “very useful,” while 23.3% considered it “moderately useful.” Only 2.3% described it as “indispensable,” and 25.6% expressed scepticism (ratings 1–2). The prevailing view positions AI as valuable yet non-essential, reinforcing the idea of a co-pilot model in which automation supports but does not replace interpreter agency. The gap between favourable utility ratings and low subscription rates points to constraints linked to usability, training, and reliability, and indicates that, although AI has moved beyond the periphery of practice, it has not yet become central.

Regarding complementarity, most interpreters endorsed AI as an adjunct to traditional preparation. When asked if AI complements document retrieval, reading, and summarisation, 37.2% agreed and 30.2% strongly agreed; smaller proportions disagreed (11.6% “strongly disagree,” 7% “disagree”), while 14% were neutral. This distribution aligns with the view that AI strengthens established routines — especially in terminology retrieval, document filtering,

and thematic anticipation — without supplanting foundational linguistic and extralinguistic preparation.

Patterns of frequency indicate moderate engagement. While 38.5% reported frequent use and 20.9% sometimes use AI, only 7% stated that they always do so. In contrast, 14% reported rarely using AI and 18.6% never. Overall, 66.3% employ AI at least occasionally, with 45.5% doing so regularly. The presence of rare and non-users suggests persistent barriers related to training, trust, and assignment relevance, and underscores the need for targeted digital literacy.

Task profiles reveal selective adoption. Terminology search is the most common application (26 yes, 9 no), closely followed by glossary creation (24 yes, 11 no). Uses are lower for documentary search (15 yes, 19 no, 1 unsure), summarisation (16 yes, 17 no, 2 unsure), terminology extraction from specialised texts (18 yes, 16 no, 1 unsure), and the generation of key questions or parallel materials for sight translation (15 yes, 20 no). The concentration in terminological tasks, which are time-intensive and cognitively demanding, suggests that interpreters employ AI as an informational scaffold while reserving activities that require interpretive judgement — context-sensitive summarisation, pragmatic reformulation, and creative rehearsal — for more cautious, human-led control.

Evaluations of reliability, efficiency, and quality are mixed. Five respondents perceived a significant improvement in preparation quality and eighteen a moderate one; eight reported no effect and four were unsure. Time assessments are more favourable: eleven reported a significant reduction and seventeen a moderate one, although one saw no change, two noted increases, and four were unsure. Qualitative comments clarify these ambivalences. Interpreters value speed and breadth but criticise generic, decontextualised outputs that demand manual correction. The oft-cited trade-off — “time saved in search is lost in correction” — highlights a shift in the interpreter’s role from sole author to editor–curator of machine-generated content. Concerns recur around confidentiality, opacity of sources, and potential skill erosion through overreliance.

Despite reservations, a hybrid approach predominates: twenty-nine respondents prefer combining traditional methods with AI, one relies exclusively on AI, and five prefer manual strategies. Calls for training in

prompt optimisation and error detection were frequent, signalling a need to move beyond operational familiarity toward interpretive discernment. In sum, AI is appreciated for streamlining repetitive components of preparation but remains insufficient to satisfy the epistemic and discursive demands of complex assignments. Its appropriate status is that of strategic assistance, providing cognitive scaffolding and information filtering while core interpretive tasks continue to rely on human expertise, ethical judgement, and domain knowledge.

Open-ended reflections on the future confirm this balanced stance. Time efficiency, rapid access to terminology, and support with last-minute documentation were the most salient benefits. Respondents valued automation of repetitive tasks such as preliminary glossary building or initial topic exploration, which frees cognitive resources for higher-order processes. Yet they warned against superficial engagement and the “false sense of readiness” produced by uncritical reliance on AI-generated outputs. Confidentiality was a central concern, especially regarding client materials and institutional contexts where data ethics and trust are paramount. Participants also mentioned environmental costs, noting the energy intensity of generative AI. Across responses, the need for training remained a constant: interpreters asked for guidance on effective prompting, validation protocols, and the integration of AI within personal preparation styles. Overall, the outlook is cautiously optimistic: AI should operate within a hybrid, ethically governed model in which interpreters remain central to their own preparation.

6 Interpreting students: pre-/post-training analysis

The student cohort offers a forward-looking perspective on emerging professional practices. Baseline data indicate moderate familiarity with AI and chatbots but limited operational knowledge of Microsoft Copilot. Expectations were generally more positive for document analysis than for glossary creation or sight translation materials, reflecting caution about tasks demanding

deeper conceptualisation. Broader attitudes revealed interest in professional applications alongside substantial concerns about automation, deskilling, and ethics (privacy and opacity). In short, students recognise AI's relevance yet approach it critically, which underlines the need for structured pedagogical integration combining technical fluency with ethical and professional awareness.

The quality of prompting proved decisive for the effectiveness of AI-supported preparation. Approximately 34.5% of students produced high-quality prompts: instructions were clear, context-rich, and strategically sequenced to elicit terminology extraction, glossary building, concept explanation, and practice materials (including cloze and sight translation texts). A larger group (41.4%) produced medium-quality prompts — functionally aligned but under-specified in scope, output format, or organising criteria — yielding useful yet limited outputs. Finally, 24.1% generated low-quality prompts characterised by vagueness and lack of structure, resulting in generic or poorly adapted materials. These patterns demonstrate that AI's pedagogical yield depends less on tool availability than on prompt literacy as a cognitive and strategic competence.

Post-intervention evaluations of Copilot were moderately positive. The mean satisfaction score was 3.37/5; the majority rated the tool between three and four stars. Students praised speed in terminology extraction and organisational support, but they frequently mentioned the need for highly specific instructions, occasional inaccuracies (e.g., confusion between “Life Cycle Analysis” and “Life Cycle Assessment”), and limited adaptability to nuanced requests or export needs. Most agreed that Copilot improved preparation, although 31% emphasised ongoing review and correction. Technology Acceptance Model indicators pointed to favourable perceptions of usefulness, ease of learning, and intentions to continue using AI, tempered by reservations about reliability and user experience. The preference for using Copilot over manual preparation was clear (72.4%), justified by time efficiency and better organisation, whereas manual preparation was preferred by 13.8% for its contribution to deeper cognitive engagement and retention. The remaining students articulated a dual effect: AI accelerates terminology and organisation but can dampen the active processing required for higher-

order competence. Difficulties reported by roughly half the group centred on the precision of prompts, limited responsiveness to nuanced tasks, and absent export functions — barriers that reinforce the need for explicit training in prompt design and quality control.

Overall, the student data align with the hybrid preparation model: AI supplements but does not replace human cognitive work. Effective integration requires guidance on prompting, verification, and the calibration of automation so that efficiency gains do not compromise depth of learning. The results thus highlight the centrality of pedagogical framing: without a model that fosters critical engagement, AI use risks remaining superficial, whereas guided exposure can transform the tool into a vehicle for metacognitive insight and disciplined practice.

7 Discussion

Findings across both studies converge on a consistent picture: AI has become an established component of pre-task preparation, valued primarily for its ability to accelerate terminology work, filter materials, and organise inputs. Professional interpreters and students alike recognise these gains but reserve judgement — and control — over tasks that require interpretive modelling, register sensitivity, and pragmatic judgement. The result is a selective, interpreter-driven integration governed by task type: automation is deployed where it reduces repetitive cognitive load; human expertise predominates where discourse construction and ethical responsibility are at stake.

The ambivalences recorded in quality and time assessments illuminate a structural shift in the interpreter's role. AI does not eliminate preparatory labour so much as reallocate it: less time is spent on initial search and extraction, while more attention is devoted to validation, calibration, and editorial decision-making. This reallocation explains why perceived time savings coexist with reports that post-editing can offset gains. It also clarifies why perceptions of quality improve only moderately: raw outputs accelerate access, but do not replace the deep structuring that interpreters achieve through close reading and

rehearsal. The editorial role thus becomes integral to professional preparation, as does explicit training in prompt optimisation and error detection.

Ethical concerns — data confidentiality, source opacity, algorithmic bias, and environmental impact — appear not as peripheral anxieties but as recurrent conditions of responsible use. Respondents emphasised that institutional frameworks and professional codes must be respected, particularly in high-stakes settings. The prudent response is not abstention but governance: privacy-preserving practices, transparent data handling, and critical evaluation of outputs should accompany any pedagogical or professional integration of AI.

Finally, the student intervention sheds light on how educational design shapes adoption. Gains in efficiency and organisation emerged when prompting was structured and reflective. Where prompting lacked specificity, outputs were superficial and learners reported friction. This gradient underscores the importance of targeted instruction and assessment that reward not only task outcomes but also process quality — how students formulate requests, justify selections, and validate results. In this respect, AI becomes a means to render tacit strategies observable and teachable.

8 Conclusions and future directions

Across professional and educational contexts, AI demonstrates a clear capacity to enhance pre-task preparation by expediting terminology management, information retrieval, and the organisation of materials. Among professional interpreters, adoption is pragmatic and selective: automation is welcomed for routine, time-consuming tasks, yet human oversight remains indispensable to ensure terminological accuracy, contextual appropriateness, and discursive coherence. Concerns about confidentiality, bias, and potential deskilling temper enthusiasm and confirm that interpreters must preserve epistemic and ethical agency within hybrid preparation ecosystems.

From an educational standpoint, guided exposure to AI tools reshapes perceptions and practices. Students reported efficiency gains and improved lexical organisation, while also identifying limitations in semantic depth,

customisation, and the sensitivity of outputs to nuanced instructions. Post-intervention attitudes reflect cautious acceptance: willingness to continue using AI depends on explicit training in prompt design, verification routines, and ethical awareness. These findings support a pedagogical model in which AI is integrated as a scaffold subject to critical governance rather than as an unexamined shortcut.

In operational terms, three implications follow. First, AI literacy should be systematically incorporated into interpreter training, with emphasis on prompt engineering, output validation, error detection, and reflective documentation of decisions. Second, hybrid preparation models should be promoted explicitly so that automation complements, rather than displaces, the deep cognitive work of analysis, thematic structuring, and rehearsal. Third, ethical training must address data privacy, algorithmic transparency, inclusivity, and environmental impact as core — not auxiliary — dimensions of professional formation.

Future research should pursue longitudinal designs to track how interpreters' relationships with AI evolve across career stages and domains of specialisation, and to examine how varying doses of automation affect retention, register control, and performance under pressure. Comparative analyses by language pair and technical field would clarify heterogeneity in impact and inform tailored pedagogical interventions. Advancing along these lines will allow the profession to realise AI's promise — speed, scale, organisation — without compromising the interpretive agency that defines interpreting as an expert, ethical, and communicative practice.

References

- ÁLVAREZ PÉREZ, B. (2024): *The Pre-Task Preparation Process of the Interpreter: A Practical Application Proposal for Glossary Development in Translation and Interpreting*. Universidad de Las Palmas de Gran Canaria. <https://accedacris.ulpgc.es/jspui/handle/10553/129017> (04.12.2025).
- ANTONOVA, A. M. (2023): "Training the Augmented Interpreter Today." *Lomonosov Linguistics and Intercultural Communication Journal* 26.4: 25–37. DOI: 10.55959/MSU-2074-1588-19-26-4-2 (13.1.2025).

- BHATTACHARYA, D. (2025): "Teachers' Competencies in Effective Technology-enabled Assessment for Holistic Development of Students." *Indian Journal of Educational Training* 7.1: 83–104. <https://journals.ncert.gov.in/IJET/article/view/843> (13.1.2025).
- BRAUN, V. / CLARKE, V. (2006): "Using thematic analysis in psychology." *Qualitative Research in Psychology* 3.2: 77–101. DOI: 10.1191/1478088706qp063oa (13.1.2025).
- BUDIHARJO, B. (2024): "Artificial Intelligence in Translation: The Menace, Promise, and Response to Technology and Superseded Practice." *Proceedings of the Third International Conference on Communication, Language, Literature, and Culture, Advances in Social Science, Education and Humanities Research* (Vol. 883). Eds Zita Rarastesia / Radhika Gajjala / Hishamundi Isam / Djatmika Djatmika. Atlantis Press. 681–694. DOI: 10.2991/978-2-38476-321-4_52.
- CHEN, S. / KRUGER, J. L. (2024): "A Computer-assisted Consecutive Interpreting Workflow: Training and Evaluation." *The Interpreter and Translator Trainer* 18.3: 380–399. DOI: 10.1080/1750399X.2024.2373553.
- CORPAS PASTOR, G. / DEFRANCQ, B. (2023): *Interpreting Technologies – Current and Future trends*. Amsterdam: Benjamins.
- CRESWELL, J. W. / PLANO CLARK, V. L. (2018): *Designing and Conducting Mixed Methods Research* (3rd ed.). Thousand Oaks, California: SAGE Publications.
- DEFRANCQ, B. / CORPAS PASTOR, G. (2023): "Introduction." *Corpas Pastor / Defrancq* (2023): 1–5.
- DOWNIE, J. (2020): *Interpreters vs Machines: Can Interpreters Survive in an AI-dominated World?* London and New York: Routledge.
- FAN, D. C. (2024): "Conference Interpreters' Technology Readiness and Perception of Digital Technologies." *Interpreting* 26: 178–200. DOI: 10.1075/intp.00110.fan.
- FANTINUOLI, C. (2023): "Towards AI-enhanced Computer-assisted Interpreting." *Corpas Pastor / Defrancq* (2023): 46–71.
- FANTINUOLI, C. / MONTECCHIO, M. (2022): "Defining Maximum Acceptable Latency of AI-enhanced CAI Tools." <https://arxiv.org/pdf/2201.02792> (12.11.2025).
- GILE, D. (2009): *Basic Concepts and Models for Interpreter and Translator Training* (Rev. ed.). Amsterdam: Benjamins. DOI: 10.1075/btl.8.
- GUO, M. / HAN, L. / TEIXEIRA ANACLETO, M. (2023): "Computer-assisted Interpreting Tools: Status quo and Future Trends." *Theory and Practice in Language Studies* 13.1: 89–99. DOI: 10.17507/tpls.1301.11.

- KALINA, S. (2015): "Preparation." *Routledge Encyclopedia of Interpreting Studies*. Ed. F. Pöchhacker. London and New York: Routledge. 318–320.
- LIU, J. (2022): "The Impact of Technologies on Interpreting: An Interpreter and Trainer's Perspective." *International Journal of Chinese and English Translation & Interpreting* 1: 1–8. DOI: 10.56395/ijceti.v1i1.14.
- LIU, K. / CHEUNG, A. K. (2023): "Translation and Interpreting in the Age of COVID-19: Challenges and Opportunities." *Translation and Interpreting in the Age of COVID-19*. Eds. K. Liu / A. K. F. Cheung. Singapore: Springer. 1–10.
- MELLINGER, C. D. / HANSON, T. A. (2018): "Interpreter Traits and the Relationship with Technology and Visibility." *Translation and Interpreting Studies* 13: 366–392. DOI: 10.1075/tis.00021.mel.
- SUN, L. (2025): "Transforming Business Interpretation Education with AI: Perspectives from Instructors and Learners." *Education and Information Technologies* 30: 11827–11861. DOI: 10.1007/s10639-024-13270-1.
- XU, R. (2018): "Corpus-based Terminological Preparation for Simultaneous Interpreting." *Interpreting* 20.1: 29–58. DOI: 10.1075/intp.00002.xu.
- YANG, S. / MU, L. (2024): "The Influence of Pre-task Preparation on Cognitive Load in English to Chinese Consecutive Interpreting Process." *International Journal of Chinese and English Translation & Interpreting* 6: 1–19.
- ZHANG, A. (2017): "Simultaneous Interpreting (SI): The Holy Grail of Artificial Intelligence – An SI Practitioner's Perspective." *Lebende Sprachen* 62.2: 253–268. DOI: 10.1515/les-2017-0023.
- ZHANG, X. / CORPAS PASTOR, G. / ZHANG, J. (2023): "Videoconference Interpreting Goes Multimodal: Some Insights and a Tentative Proposal." *Corpas Pastor / Defrancq* (2023): 169–194.

Ist KI überzeugend? Dolmetschen in Weinverhandlungssituationen

Abstract

Business negotiations are communicative acts in which persuasion is essential to achieving communicative intent. In both import and export contexts, negotiation partners frequently rely solely on business interpreters, making it necessary for prospective interpreters to be trained for such situations. Artificial Intelligence has recently gained prominence in interpreting, particularly through chatbots such as Copilot, raising questions about the appropriateness of these tools in professional and training interpreting settings. This study analyses the suitability of Copilot in two contexts: (a) preparing a wine-import negotiation scenario (German–Spanish) for interpreter training, and (b) performing the interpretation in the same type of negotiation. The performance of human interpreters is compared with that of provided by Copilot. The initial hypothesis posits that Artificial Intelligence lacks the capacity to perform the full range of speech acts required in business settings, and that its inability to do so may hinder the persuasive function central to the business communicative framework.

Keywords Geschäftsverhandlungssituationen, KI, kommunikatives Versagen, menschliche versus maschinelle Übersetzung/Verdolmetschung, Verhandlungsdolmetschen

1 **Einleitung: Geschäftsverhandlungen und Dolmetschen – Umbruch durch KI**

Künstliche Intelligenz (KI) hat nicht nur Einzug in unseren Alltag gehalten, sondern kommt immer häufiger in Fachkontexten zur Anwendung. Im Übersetzungsbereich wird KI inzwischen nicht nur für das Übersetzen verwendet, sondern auch für Dolmetschaktivitäten (Jiménez Serrano 2025). Aber auch für andere Kommunikationskontexte ist KI gefragt. Chatbots, zum Beispiel, werden immer öfter von Vortragenden zur Erstellung der Beiträge selbst benutzt, und auch im Bildungsbereich bedienen sich Lehrende der KI in der Vorbereitung verschiedener Übungen. Ein Anwendungsbeispiel in der Dolmetscherausbildung ist die Erstellung von Gesprächssituationen zum Üben des Verhandlungsdolmetschens.

Unser Forschungsansatz geht davon aus, dass eine von KI erstellte Gesprächssituation, eben dadurch, dass diese künstlich erstellt ist, auch künstlich wirkt und dadurch die Effizienz der Kommunikation aufs Spiel gesetzt wird, das heißt das Kommunikationsziel der Verhandlung (zu überzeugen). Pimentel Siles (2017) versteht *Verhandeln* als eine grundlegende menschliche Beziehung, in der die Sprachanwendung entscheidend für den Erfolg einer Verhandlung ist, da der Prozess des Aushandelns hauptsächlich ‚überzeugendes Sprechen‘ bedeutet, also ‚anbieten und argumentieren‘, ‚verlangen und geben‘, ‚annehmen oder verweigern‘. Es ist also hypothetisch davon auszugehen, dass die Kommunikation stark beeinträchtigt werden kann, wenn von KI-Tools erstellte Gesprächssituationen später auch von einer KI übersetzt bzw. verdolmetscht werden, was voraussichtlich mit der aktuellen Tendenz zur KI-Anwendung in nahezu allen Bereichen eine durchaus mögliche Situation ist (van der Meer 2020). Die Kommunikationsintention kann in Geschäftsverhandlungskontexten zunichte gemacht werden, weil wir als Menschen stets intentionsbedingt kommunizieren und daher auch die Vorbereitung einer Verhandlung strategisch angegangen werden muss, also auch mit Vorentscheidungen bezüglich der Sprache, die wir in jeder Verhandlungsphase anwenden wollen (Pimentel Siles 2017).

Pimentel Siles (2017) betont, dass eine Verhandlung mit Fragen zu WAS, WEM, WIE VIEL, WIE, FÜR WAS und WARUM strukturiert werden soll-

te. Dies ist insofern interessant, als diese W-Fragen identisch mit den Fragestellungen für die Übersetzung in der Funktionstheorie von Nord sind (vgl. Nord 2014). Nord hält diesen funktionalistischen Ansatz mit der Antwort auf die W-Fragen für unerlässlich, wenn bei der Übersetzung eines Textes der Kommunikationszweck oder „Skopos“, wie von Reiss/Vermeer (2014) schon im Jahr 1984 benannt, erfüllt werden soll.

Bravo Molina et al. (2019) gehen in diesem Sinne davon aus, dass ohne eine effektive Kommunikation eine Verhandlung nicht möglich ist, denn nur beim Kommunizieren werden die Verhandlungsziele klar ausgedrückt und verstanden sowie auch das Verhandlungsinteresse und die Wünsche beider Verhandlungsparteien. Nur so kann es zu einer Vereinbarung kommen. Diese Forscher weisen weiter darauf hin (Bravo Molina et al. 2019), dass schon aus Aristoteles' Rhetorika klar hervorging, dass Überzeugen das Hauptziel der Kommunikation ist und so Angebot und Nachfrage allein nicht den Verhandlungsinhalt ausmachen, sondern eine auf zwei Ebenen praktizierte Sprache: eine auf der Logik basierten, in der Vorschläge und Angebote erfolgen, und eine auf Pragmatik gestützte (Semantik, Syntax und Stil), in der die Wortwahl Zeichen der eigenen Haltung ist.

Fünf Dimensionen sind bei einer Kaufverhandlung wichtig: der Gewinn, die Exklusivität, das Vertrauen und die Sicherheit, die Präsenz einer bestimmten Dringlichkeit und das Schaffen positiver Emotionen. Der Gewinn wird linguistisch mit Verben wie *verbessern*, *sparen*, *gewinnen*, *optimieren* und *steigern* zum Ausdruck gebracht. Adjektive, die geringe Verfügbarkeit, Einzigartigkeit und Personalisierung hervorheben, weisen auf Exklusivität hin. Das Vertrauen und die Sicherheit werden mit Adjektiven wie *garantiert*, *sicher* und Nomen wie *Unterstützung*, *Schutz* und Adverbien wie *verlässlich* hergestellt. Der Dringlichkeitsansatz wird mit dem befristeten Produktangebot mit Bezug auf Zeit- und Mengenangaben gefördert, zum Beispiel *zeitlich begrenztes Angebot*, *die letzten Einheiten* und *nur heute*. Und die Emotionen der Verhandlungsparteien stehen mit Adjektiven im Zusammenhang, die auf Überraschung und Wunsch-erwecken oder auf Schönheit hinweisen, wie etwa *unglaublich*, *überraschend*, *fantastisch*, *bewegend* und *wunderbar*.

Diese Dimensionen müssen also in den Kategorien, die eine Verhandlung ausmachen, beachtet werden. Die Interessen sieht man an den Ansprüchen,

Sorgen, Zielen, Gründen, Hoffnungen und Ängsten der Verhandlungspersonen. Alternative Möglichkeiten werden hervorgebracht, wenn es nicht zur Einigung zwischen den Parteien gekommen ist. Die Legitimität gründet darauf, dass die Parteien die Verhandlungssituation als fair einschätzen. Zu diesen Kategorien gehören auch die Kommunikation, die Vertrauensbeziehung und die wechselseitige Verpflichtung, und alle beinhalten objektive sowie subjektive Aspekte, wie die Notwendigkeit der Parteien, mögliche Lösungen, Gültigkeitskriterien, partnerschaftliche Kommunikationsweise, die erwünschte vertrauensvolle Beziehung und das erwünschte Endergebnis.

Zudem werden fünf Verhandlungsstile unterschieden: (1) Anpassungsstil, (2) Wettbewerbsstil, (3) Kooperationsstil, (4) Verpflichtungsstil und (5) Vermeidungsstil. Genutzte Werkzeuge für eine Verhandlung sind in diesem Kontext eine effektive Kommunikation, emotionale Intelligenz, Informationsbeschaffung, Überzeugungskraft und Optionsanalyse. In der Verhandlung kann man außerdem verschiedene Phasen unterscheiden: die Vorbereitung, die Diskussion, die Erklärung, den Vorschlag, die tatsächliche Verhandlung und den Geschäftsabschluss. Die Verhandlung kann außerdem, schließlich, auch teilbelastet (gewinnen-verlieren, *win-lose*) oder integrativ (gewinnen-gewinnen, *win-win*) sein.

2 Zur Forschung: Hypothese, Ziele, Materialien und Methoden

Wie schon in der Einleitung ausgeführt, ist eine Geschäftsverhandlung ein Kommunikationskontext, in dem der wichtigste Sprechakt, neben *Verhandeln*, *Überzeugen* und *Vereinbaren*, auch *Beschenken* (vor allem im übertragenen Sinne) ist. Aus diesem Grund ist die genaue Wortwahl zum Erlangen der erwünschten Effekte wichtig und Dolmetschende müssen diese in ihren Arbeitssprachen in der Ausgangsrede richtig erkennen und adäquat in die Zielsprache übertragen.

In Import- und Exportverhandlungen wird üblicherweise Dialogdolmetschen von Seiten der Kunden beansprucht. Der Dolmetscher ist in dieser

Verhandlungssituation oft ein weiteres wichtiges Element des geschäftlichen Erfolgs. Unsere Hypothese ist, dass je mehr KI eingesetzt wird, desto unnatürlicher wird die Sprache der Verhandlungssituation und somit die kommunikative Effizienz. Diese Hypothese geht davon aus, dass KI einerseits für die Vorbereitung einer Übung für Dialogdolmetschen im Rahmen einer Dolmetscherausbildung nicht geeignet ist, da diese künstlich erstellt ist und auch künstlich wirkt und dadurch für die Übung einer geeigneten Effizienz der Kommunikation auch nicht nützlich ist.

Andererseits wird auch davon ausgegangen, dass KI bei der Verdolmetschung einer Verhandlungssituation ebenso nicht so effektiv ist wie ein Dolmetscher.

Um diesen Hypothesen nachzugehen, stellen wir folgende Forschungsziele auf:

- A) Vergleich einer mit KI erstellten mit einer von Menschen durchgeführten Verhandlungssituation
- B) Vergleich der Verdolmetschung einer menschlich und maschinell (mit KI) erstellten Situation
- C) Vergleich, wie Mensch und Maschine eine von einem Menschen vorbereitete Verhandlungssituation verdolmetscht

Forschungsgrundlage der Studie bilden: (a) eine maschinell (Copilot¹) erstellte Verhandlungssituation (s. Anhang 1); (b) eine von Menschen vorbereitete² Verhandlungssituation (s. Anhang 1); (c) eine maschinell (Copilot) durchgeführte Verdolmetschung der maschinell erstellten und der menschlichen Ver-

.....

- 1 Die Autorinnen haben sich für diese KI für die Studie entschieden, weil es sich bei Copilot, KI-Assistent von Microsoft, um einen privaten und sicheren Chat-Dienst handelt, der gewährleistet, dass die Daten innerhalb des Unternehmens bleiben und nicht an Dritte weitergegeben werden.
- 2 Von einer der Autorinnen dieses Beitrags, Macarena Pradas, für ihren Dolmetschunterricht (Deutsch-Spanisch) an der Universität Granada.

handlungssituation (s. Anhang 2) und (d) eine menschliche Verdolmetschung der maschinellen und menschlichen³ Verhandlungssituation (s. Anhang 3).⁴

Die menschliche Verhandlungssituation geht von einer integrativen Verhandlung (gewinnen–gewinnen) aus, für die fiktive Gesprächspartner (Spanisch und Deutsch) und ein erfundenes Geschäftsziel erstellt wurden (vgl. Tabelle 1). Diese fiktiven Angaben wurden auch Copilot zur Erstellung der Verhandlungssituation als Prompt vorgegeben. Die Situation entspricht also einer fiktiven, aber realitätsnahen Verhandlungssituation⁵, damit sich Dolmetschstudierende realitätsnah auf die Situation vorbereiten können. Gesprächspartner sind Herr Winzer (Wiener Winzer in Österreich) und Frau Racimo (Vertreterin einer Weinkellerei in Katalonien, Spanien). Ausgemachter Ort des Treffens ist die Weinmesse VieVinum in Wien (Österreich) und das Verhandlungsziel ist die Exportvertretung für die österreichische Weinkellerei durch das spanische Unternehmen.

Abbildung 1 enthält den Prompt zur Erstellung der Situation durch Copilot (ursprünglich auf Spanisch):

Guten Tag, Copilot. Kannst du mir bitte ein Gespräch erstellen, das die Verhandlung zwischen Frau Racimo (Vertreterin einer katalanischen Weinkellerei) und Herrn Winzer, Wiener Winzer, der sein Wein exportieren will, darstellt?

Die Entwicklung der Verhandlung ist wie folgt: Beide Gesprächspartner treffen sich auf der Weinmesse, Vievinum in Wien. Sie haben ausgemacht, sich dort zu treffen. Im Rahmen dieses Treffens stellen beide ihre Produkte

.....

3 Wie im Falle der Verhandlungssituation (Fußnote 2), wurde diese Übersetzung/Verdolmetschung von Macarena Pradas für den Dolmetschunterricht in Granada vor Jahren durchgenommen.

4 Das Korpus (Anhang 1–3) ist einsehbar unter folgendem Link: https://drive.google.com/drive/folders/1ybaoiBoecXcstcMZGc1Cuw_BZV8_m9FW?usp=sharing (28.01.2026).

5 Zur Erfüllung dieses Kriteriums ist die Situation komplett erfunden, aber der Treffort, VieVinumMesse in Wien, die Orte (z. B. Heuriger Hengl-Haselbrunner) existieren in der Realität. Sollten die Namen der Gesprächspartner oder sonstige andere Details der Wirklichkeit entsprechen, ist das nur Zufall.

vor und beschreiben deren Eigenschaften und auch die Unternehmen (katalanischer Weinvertrieb und Wiener Weinkellerei).

Erste Szene: Beide Gesprächspartner stellen sich vor und sprechen über ihre Produkte und Dienstleistungen.

Szene 2: Beide besprechen die Importbedingungen von Seiten des katalanischen Weinvertriebs.

Szene 3: Schließlich lädt der Wiener Winzer die katalonische Vertreterin zum Essen in einen typischen Heurigen in Wien ein (Hengl-Haselbrunner). Die Verhandlung soll in einem für ein Luxusgut wie Wein entsprechendem Sprachstil und -register Fachsprache zum Thema beinhalten (Anbau, Weinlese, Weinherstellung und Weinverkostung). Außerdem soll in der Situation auch auf den Treffpunkt der Verhandlung, die Hofburg in Wien, wo die Weinmesse VieVinum stattfindet, eingegangen werden sowie auf die Speisekarte im Heurigen Hengl-Haselbrunner. Weiter soll beschrieben werden, wie man diesen Heurigen am besten erreichen kann (Transportmöglichkeiten). Der Wiener Winzer hat Deutsch als Muttersprache und die katalanische Vertreterin spricht Spanisch.

Abb. 1: Prompt für Copilot zur Erstellung der Importverhandlungssituation

Für den Vergleich der menschlichen und maschinellen Verhandlungssituation sind wir wie folgt vorgegangen: In Bezug auf die Verhandlung haben wir den Verhandlungsstil untersucht, die zur Verhandlung nötigen Instrumente (effektive Kommunikation, emotionale Intelligenz, überzeugen) und welche Phasen der Verhandlung aufgenommen werden (Vorbereitung, Diskussion, Klassifizierung, Vorschlag, tatsächliche Verhandlung und Geschäftsabschluss) (s. Anhang 1).

Anschließend wurde festgestellt, inwieweit die fünf für eine Verhandlung wichtigen Dimensionen respektiert werden (s. Abschnitt 1). Dafür wurde das Vorkommen jeder Kategorie in der menschlichen und in der maschinellen Verhandlung untersucht und die Anzahl der für die Verhandlung zur Dimensionserfüllung verwendeten Verben berechnet. Auch die Verwendung von Namen und explizite und implizite Pronomen wurden in der jeweiligen Verhandlungssituation markiert (kursiv in Anhang 1), um die Platzierung im Satz besser

beobachten zu können. Die Anwendung phatischer Partikel, die das Interesse für ein Gespräch unterstreichen, wurden hervorgehoben (s. Anhang 1).

Für den Vergleich zwischen menschlicher und maschineller (Copilot) Übersetzung/Verdolmetschung⁶ haben wir die sieben Parameter⁷ von Mos-sop (2014) zur Qualitätsüberprüfung der vorgeschlagenen Kriterien benutzt, da diese aus unserer Sicht alle der vier Übersetzungsdimensionen von Mos-sop (2014) (Übertragung/Transfer, Inhalt, Sprache und Präsentation) gut abdecken. Rechtschreibung und Zeichensetzung zum Beispiel sind hier nicht essenziell, da diese Situationen eher für den Einsatz im mündlichen Kontext gedacht sind (Qualität der Übersetzung bei einer Tonanwendung zum Einsatz als Verdolmetschung in einer Verhandlungssituation). Idiomatik und Phraseologie haben wir in einer Kategorie erfasst, da beide Parameter sich manchmal überlappen.

3 Forschungsergebnisse und Diskussion

3.1 Vergleich der erstellten Verhandlungssituationen (Mensch und Copilot)

3.1.1 Verben

Beim Vergleich beider Verhandlungssituationen fällt sofort auf, dass die maschinelle Version viel kürzer ausfällt als die menschliche Version. Die Wörteranzahl in der menschlichen Situation unterteilt sich in 967 Wörter in der ersten Situationsszene, 400 in der zweiten und 193 in der dritten, wobei die deutschen Teile bei den ersten zwei Szenen kürzer ausfallen als die spanischen (294 Szene 1, 190 Szene 2), aber viel länger in der dritten Szene, in der 166 der 193 Wörter vom deutschen Teil eingenommen werden. Das ist sehr sinnvoll,

.....
6 Copilot hat eine Funktion zur Aussprache der schriftlichen Übersetzung und konnte daher als Dolmetsch-Tool verwendet werden.

7 Genauigkeit, Vollständigkeit, Terminologie, Phraseologie, Idiomatik, Fakten, Sprachverwendung

wenn man bedenkt, dass der letzte Teil (Szene 3) die gefundene Vereinbarung betrifft und der deutsche Teil die noch offene.

Bei der maschinellen Verhandlungssituation fällt im Gegensatz zur menschlichen Situation auf, dass auch die Unterteilung anders ist. In der ersten Szene entsprechen dem deutschen Teil 215 der insgesamt 322 Wörter, also fast zwei Drittel des Gesprächs. In der zweiten Szene haben zu fast 50 % beide Teile das Wort (73 der 146 Wörter fallen auf den deutschen Teil) und in der dritten Szene schließlich kommt die spanische Partei kaum zu Wort mit nur 9 der 89 Wörter.

Dieses Ergebnis wirkt sich auch auf die Phasen der Verhandlung aus, da alle viel kürzer ausfallen und auf manche wird in Copilots Version (Diskussion in Szene 2 und Schluss in Szene 3, siehe Anhang 1) sogar kaum eingegangen. Wie aus der genauen Untersuchung der Verwendung lexikalischer Kategorien (Tabellen 1 bis 7) entnommen werden kann, sind in der Copilot-Version weniger Elemente für eine effektive Kommunikation und zur Anregung der emotionalen Intelligenz enthalten und das kann somit zu einer reduzierten Überzeugungskraft führen.

Beim Vergleich der beiden Verhandlungssituationen (Mensch und Copilot) konnte man folgende Verwendung der Verben in den deutschen Teilen der menschlichen Verhandlungssituation feststellen (vgl. Tabelle 1).

Tab. 1: Verb-Verwendung in den deutschen Teilen der vom Menschen erstellten Verhandlungssituation

Szene 1	Szene 2	Szene 3
erfüllt werden müssen, erwartet werden, exportieren, für angepasst halten, glauben, gut beschreiben, hierherkommen, importie- ren, in Vertretung haben, interessiert sein (3x), kommen, können, etwas machen können, mich freuen, mitbringen, möchten, Recht haben, sehen (2x), sein (8x), sich bedanken, sich um etwas handeln, sich zurückbelaufen auf, vertreten, wissen (3x), X kann nicht Y ausfallen, zur Kost empfehlen, zurechtfinden, zurück- geben, zusagen	ausgezeichnet werden, die Wahrheit sein, erhalten, haben, sein, sich bemühen um, sich verteilen auf, umfassen, verarbeiten, vorbereiten, wachsen in, Zertifikat für etwas sein	anbieten, ausschenken, bezeichnen, dabei helfen, erreichen, gerne einladen, gerne wollen, können, liegen in, sein (2x), sich handeln um, sich treffen, sich um höchste Qualität bemühen, um Reservierung bitten
32 Verben	12 Verben	15 Verben
14,29 % (von 294 Wörtern)	6,32 % (von 190 Wörtern)	9,04 % (von 166 Wörtern)

Die Wahl der deutschen und der spanischen Verben fördert die Dimensionen Vertrauen und Sicherheit (z. B. *sich zurückbelaufen auf*, *ausgezeichnet werden*; *pasar a la cata*, *satisfacer*), die Präsenz einer bestimmten Dringlichkeit (z. B. *in Vertretung haben*; *seleccionar*) und das Schaffen positiver Emotionen (z. B. *erwartet werden*, *gut beschreiben*, *mich freuen*; *sorprender*).

Die in den spanischen Teilen der menschlichen Verhandlungssituation identifizierten Verben sind Tabelle 2 zu entnehmen.

Tab. 2: Verb-Verwendung in den spanischen Teilen der vom Menschen erstellten Verhandlungssituation

Szene 1	Szene 2	Szene 3
acompañar, gustar (2x), ampliar (2x), importar (2x), buscar, incorporar, caracterizarse por, interesar, coincidir con, ir, competir, llegar, comprobar, llevar, considerar, montar, convertirse en, motivarse, deber, observar, dirigir, ofertarse, distribuir (2x), ofrecer, empezar (3x), organizarme, encontrar, pasar a la cata, encontrarse entre, poder, envasarse, presentarse, estar en, probar, estar formado por, estar impresionada, representar (4x), satisfacer, representar algo, ser (9x), seleccionar, ser difícil, ser exigente, ser imposible, situarme, sorprender, suministrar, tener (2x), venir, ver, visitar	abusar, apreciarse, considerar, contactar, controlar, cumplir con (2), garantizar (2), llegar a un acuerdo (2), poder perfilar, ecibir reconocimiento, saber, ser (4), ser importante (2), ser imprescindible, situarse, tener (3), ver	aceptar, dirigirse mejor a, encontrarse, llamar a un taxi, ser
69 Verben	23 Verben	5 Verben
10,25 % (von 673 Wörtern)	10,95 % (von 210 Wörtern)	18,52 % (von 27 Wörtern)

In Tabelle 3 kann man das Aufkommen von Verben in den deutschen Teilen der maschinellen Verhandlungssituation (Copilot) sehen und in Tabelle 4 die Verwendung in den spanischen Teilen.

Tab. 3: Verwendung der Verben in den deutschen Teilen der maschinellen Verhandlungssituation

Szene 1	Szene 2	Szene 3
<i>anbauen, arbeiten, ausbauen, bewahren, enthalten, erfolgen (3x), extrahieren, filtrieren, haben, in etwas bringen, liegen, präsentieren, reifen, sein (3x), setzen auf, sich auszeichnen, sich freuen, treffen, verzichten, wachsen</i>	<i>bereitstellen, betragen, geben (2x), können (2x), liegen bei, nominieren, sein (3x)</i>	<i>anbieten, einladen, weitersprechen</i>
24 Verben	11 Verben	3 Verben
11,16 % (von 215 Wörtern)	15,07 % (von 73 Wörtern)	3,75 % (von 80 Wörtern)

Tab. 4: Verwendung der Verben in den spanischen Teilen der maschinellen Verhandlungssituation

Szene 1	Szene 2	Szene 3
<i>conocer, dar, emplear, especializarse en, hablar, permitir, presentar, ser, trabajar, utilizar, valorar (2x)</i>	<i>augurar, destacar, establecer, obtener, poder, realizar, tener</i>	<i>servir, llegar</i>
12 Verben	7 Verben	2 Verben
11,21 % (von 107 Wörtern)	9,59 % (von 73 Wörtern)	22,22 % (von 9 Wörtern)

Bei der Wahl der deutschen und spanischen Verben können Verben wie *sich auszeichnen*, *nominieren* und *especializarse en* Vertrauen und Sicherheit her-

stellen, aber Verben zum Schaffen positiver Emotionen kommen kaum vor (*augurar, valorar*). Auffallend ist, dass die Verben eher denen entsprechen, die bei Fachgebieten, z. B. bei einer technischen Beschreibung, zur Anwendung kommen (*anbauen, extrahieren, filtrieren; emplear, especializarse*).

3.1.2 Adjektive und Adverbien

Tabelle 5 enthält die Anwendung von Adjektiven und Adverbien in beiden Verhandlungssituationen (Mensch und Copilot) in beiden Sprachen (Deutsch und Spanisch).

Tab. 5: Adjektive und Adverbien in den spanischen und deutschen Teilen der maschinell und der menschlich erstellten Verhandlungssituation

Adjektive und Adverbien der Verhandlungssituation (Mensch)	Adjektive und Adverbien der Verhandlungssituation (Copilot)
Deutsch	
<i>außerordentlich, pünktlich, gut, herzlichst, gern (3), sehr (2), bester, erstklassiger, verständlich, bekannte, neu, besser, wirklich einzigartig, qualitative (2), hochwertige, sonnenverwöhnte, ressourcenschonend, regelmäßig, höchste (2), wichtigste, kleiner, beschränktem, reichhaltig, traditionell</i>	<i>hervorragend, sehr interessant, selektiv, kalkhaltig, gut, ideal, ausschließlich, kleinen, unversehrt, spontan, kontrollierter, optimal, neu, gebraucht, grob, vegan, minimal, beste, herzlich</i>
Spanisch	
<i>bien, gratamente, magnifico, impresionante, brillante, efectivamente (3), más importante, nuevos, consolidados, deseados, establecidas, fundamental, en exclusiva, amplio, nuevas (4), mayor (3), reconocidas, pequeños, conocidos, realmente difícil, normalmente, menos considerado, suaves, equilibrada, minima, importante (2), básico, mejores, altas, imprescindible, estupendo, encantada</i>	<i>impresionante, majestuoso, respetuoso con el medio ambiente, artesanal, sostenible, vertical, equilibrada, dulce, sedosa, excelente, definitiva</i>

3.1.3 Pronomen, Namen und emphatische Partikel

In der menschlichen Verhandlungssituation beginnt fast jeder Sprechbeitrag bejahend oder mit einer phatischen Partikel an (unterstrichen in Anhang 1), während in der Copilot-Version die Bejahung nur zweimal insgesamt auftritt und als emphatische Partikel: *natürlich, sehr interessant*.

Obwohl in beiden Verhandlungssituationen Partizipien kaum vorkommen, was für eine Gesprächssituation, also eine mündliche und nicht schriftliche Sprachhandlung, üblich ist, sticht bei der Analyse der Adjektive klar hervor, dass vor allem im Falle der Anwendung auf Spanisch in der Verhandlungssituation von Copilot die Adjektive eher auf die Eigenschaften des Produkts und nicht so sehr auf Vertrauen (*consolidado, establecida, reconocida, conocido* u. a.), Exklusivität (z. B. *en exclusiva, magnífico, außerordentlich* und *einzigartig*) und positive Emotionen wie im Falle der menschlichen Verhandlungssituation (*encantadas, herzlichst, gern* usw.) ausgelegt sind.

Bei der Analyse der Adjektive kann man auch beobachten, dass die Copilot-Verhandlungssituation technische Beschreibungen bevorzugt (*kalkhaltig, kontrolliert, grob, unversehrt, selektiv, respetuosa con el medio ambiente*). Diese Beschreibung wäre eher für den direkten Großhandelsverkauf an Einzelhändler geeignet, nicht aber auf einer Messe, in der die Überzeugungskraft, damit es zu einer Vereinbarung kommt, eher mit emotionalen (Schönheit, Exklusivität usw.) und nicht mit technischen Argumenten erlangt wird.

Weiter ist auffällig (s. Anhang 1), dass in der maschinellen Situation die spanische Vertreterin anfänglich das Wort nimmt, obwohl sie ja von Herrn Winzer eingeladen worden ist und deshalb auch aufgesucht wird. Weiter wirkt es auch absonderlich, dass schon gleich im ersten Satz, bevor sich beide Gesprächspartner vorgestellt haben, sofort auf andere Aspekte (Schönheit der Platzierung, Hofburg) eingegangen wird.

Bezüglich der Verwendung von Namen und Pronomen kann man feststellen, dass in der menschlichen Situation diese von beiden Gesprächspartnern öfter im Satz und auch häufig an erster Stelle des Satzes angewendet wurden (s. Anhang 1 kursiv). Auch der gegenseitige Bezug trat häufig auf, wie aus dem folgenden Beispiel hervorgeht:

*W: Ja. Das halte **ich** für sehr angepasst. **Ich** bin sehr an dem Export **unseres** Weins interessiert, aber **ich** weiß, dass harte Bedingungen für die Einführung in Spanien erfüllt werden müssen. **Ich** weiß auch, dass **Sie** die beste Wein-Import- und -Export-Firma in Spanien sind und deshalb war **ich** so sehr am Kontakt mit **Ihnen** interessiert. Übrigens für die Verkostung des Weines habe **ich** einige Weingläser aus der Gläserstation mitgebracht. So können **wir** diese später wieder zurückgeben.*

Wenn man die Ergebnisse der Verwendung aller lexikalischer Kategorien zusammennimmt (Tabellen 1 bis 5), geht hervor, dass die Copilot-Version weniger Elemente der fünf Dimensionen für eine effektive Kommunikation und zur Anregung der emotionalen Intelligenz erfüllt und dass dies somit zu einer reduzierteren Überzeugungskraft und eher einer Nichterfüllung des Kommunikationsziels, eine Importvereinbarung zu treffen, führen kann.

3.2 Vergleich der maschinellen und der menschlichen Verdolmetschung der menschlichen Verhandlungssituation

In Tabelle 6 werden die Ergebnisse zur Qualität beider Übersetzungen (Mensch und Copilot) der menschlichen Verhandlungssituation analysiert. In der maschinellen Übersetzung konnten in allen Kategorien Fehler verzeichnet werden.

Tab. 6: Vergleich der menschlichen und der maschinellen Übersetzungen (der menschlichen Verhandlungssituation) in Bezug auf Fehler

Menschliche Verhandlungssituation						
	Richtigkeit (R)	Vollständigkeit (V)	Terminologie (T)	Phraseologie/Idiomatik (PI)	Fakten (F)	Sprachverwendung (S)
Menschliche Übersetzung						
Maschinelle Übersetzung	R1	V1 / V2	T 1 (7x)	PI1 / PI 2 / PI3 / PI4	F1	S1 / S2/ S3 / S4 / S5

Durch die falsche Verwendung der Terminologie (T1) (sieben Mal im Gespräch) wurde das Kriterium Korrektheit nicht immer erfüllt. Statt *Weinkellerei* wurde *bodega* aus dem Spanischen mit *Weingut* ins Deutsche übertragen. Da es in der Verhandlungssituation genau darum geht, eine *Weinkellerei* in Vertretung zu nehmen, wurde die Genauigkeit damit nicht erlangt. *Weingut* bezieht sich auf das Anbaugebiet der Trauben und *Weinkellerei* auf das Gebäude, in dem der Wein hergestellt/gelagert wird. Da bei *Weingut* der Fokus beim Anbau liegt, ist es auch verständlich, dass in der maschinellen Verhandlungssituation tiefer darauf eingegangen wurde als in der menschlichen, wo die Qualität des Weines im Mittelpunkt stand. Auch wurde *herzlichst dafür bedanken* mit *agradecerle sinceramente* übersetzt und sollte besser *agradecerle de corazón* nutzen, da sonst mehr Distanz entsteht, die für diese Kommunikationssituation nicht angestrebt wird.

Die Vollständigkeit (V) wurde zweimal nicht erreicht. *Ich möchte mich erstmal(s) ganz herzlichst dafür bedanken, dass sie hierhergekommen sind* wurde mit *Quiero agradecerle sinceramente que haya venido* wiedergegeben. *Erstmals und hierher* wurden ausgelassen.

Auch mit Blick auf Phraseologie und Idiomatik (PI) wurden vier Fehler festgestellt, die das Sinnverständnis beeinträchtigen könnten. Bei PI1 z. B. wurde *Sie waren außerordentlich pünktlich, genau 16 Uhr* mit *Ha sido extraordinariamente puntual, exactamente a las 16 horas* übersetzt. Einerseits benutzt man im Spanischen in einer mündlichen Situation die Zeitangaben immer nur von 1 bis 12 Uhr mit der Angabe *hora*, und dann wird höchstens gesagt, ob sich diese auf vor oder nach der Mittagszeit beziehen. Auch klingt es nicht idiomatisch, kein Verb zu benutzen. Es sollte zum Beispiel besser mit *ha llegado exactamente a las cuatro de la tarde* und idiomatischer als *exactamente* ist in diesem Kontext *en punto, ha llegado a las cuatro en punto/son las cuatro en punto*. PI2 ist zwar sprachlich richtig, aber für diesen Kontext nicht passend.

Übersetzt wurde *Wir können gerne einen Rundgang durch die Hofburg machen* mit *Podemos hacer con gusto un recorrido por la Hofburg*. Im Deutschen handelt es sich um eine Einladung, im Spanischen hingegen wäre diese fast wörtliche Übersetzung eher eine Antwort auf diese Einladung, da sonst normalerweise auf irgendeine Weise, zum Beispiel mit einer Fragestellung, die Meinung dazu erwartet werden sollte. PI3 und PI4 sind nicht idiomatisch.

Bei PI3 geht es um die Formulierung: *Damit Sie Beispiele für die Qualität der Weine sehen*. Idiomatischer wäre *sehen können* oder *an Beispielen erkennen können*. PI4 entspricht dem Übersetzungsvorschlag *Wir begannen mit neuen Projekten, die heute gefestigt sind*. In diesem Fall wäre *solide* angemessener.

Ein Fehler bezüglich der Fakten kam nur einmal vor (F1). Zwar ist Burgund die richtige Übersetzung für *Borgoña*, wenn die Region gemeint ist, aber im Weinkontext verwendet man *Bourgogne*.

Die fünf falschen Sprachverwendungen stehen vor allem im Zusammenhang mit der Satzstruktur. Im Falle von S1 steht *distribuidora catalana de vinos*, wobei die Funktion des Vertriebs vor dem Standort kommen sollte und *distribuidora de vinos catalana* stehen. Bei S2 wirkt es nicht natürlich, dass man in einer mündlichen Situation die Herkunftsangabe vor der eigentlichen Personenidentifizierung angibt. Es heißt *Sie müssen Herr Winzer von der Wiener Weinkellerei sein* und sollte eher korrekt heißen *Sie müssen Herr Winzer sein* und dann *von der Wiener Weinkellerei*. In den Fällen S3, S4 und S5 handelt es sich um falsche Satzstrukturen, die das Sinnverständnis erschweren: (S3): *Um die Stände der Weingüter aufzubauen, die uns begleiten* – besser: *um die Stände der Weingüter, die uns begleiten, aufzubauen*; (S4): *damit Sie Beispiele für die Qualität der Weine sehen, die wir vertreiben* – besser: *damit Sie die Qualität der Weine, die wir vertreiben, direkt probieren*; (S5): *Wir importieren die Weingüter, die uns gefallen, mit spannenden Projekten* – besser: *Wir importieren die Weingüter mit spannenden Projekten, die uns gefallen*. Bei S6 ist zum richtigen Verständnis ein anaphorischer Hinweis wichtig (*Die grundlegenden Anforderungen nehme ich an, dass Sie erfüllen* – besser: ..., *dass Sie sie erfüllen*).

3.3 Vergleich der maschinellen und der menschlichen Verdolmetschung der maschinellen Verhandlungssituation

Tabelle 7 enthält die Anzahl der Fehler pro Überprüfungskategorie. Sind die Felder leer, sind in diesen Kategorien keine Fehler entstanden.

Tab. 7: Vergleich der menschlichen und der maschinellen Übersetzungen (der maschinellen Verhandlungssituation) in Bezug auf Fehler

Verhandlungssituation Copilot						
	Richtig- keit (R)	Vollständig- keit (V)	Termino- logie (T)	Phraseologie/ Idiomatik (PI)	Fakten (F)	Sprach- verwendung (S)
Menschliche Übersetzung						
Maschinelle Übersetzung			T1	PI1 / PI2 / PI 3 / PI 4		

Bei der Analyse dieser Fälle gilt es im Auge zu behalten, dass die maschinelle Verhandlungssituation mit 557 Wörtern insgesamt ‚quasi‘ dreimal kürzer ausfällt als die menschliche Verhandlungssituation (1560 Wörter). Trotzdem kam es zur gleichen Fehleranzahl (4) in der Kategorie Phraseologie/ Idiomatik. In einem Kontext, in dem diese Kategorie so wichtig zur Funktionserfüllung ist, wiegt ein falsches Ergebnis schwer. Auch, dass in dieser Situation der gleiche Fehler (T1) wie in der menschlichen Situation entstanden ist (s. Erklärung für T1 in Tabelle 7), gefährdet wieder einmal das Verständnis.

In Bezug auf die Phraseologie/Idiomatik wurde im Fall von PI1 *Es freut mich sehr, Sie hier zu treffen* mit *Me alegre mucho encontrarla aquí* übersetzt. *Encontrarla* würde man im Spanischen eher für das Ausdrücken von Zufall benutzen. Wenn *encontrar* verwendet wird, dann reflexiv *encontrarme aquí con usted*, um den Sinn im Deutschen richtig zu übertragen. In PI2 erscheint *vorstellen* – wenn man damit den Vertrieb meint, ist es nicht ganz geeignet. Passender scheint *präsentieren*. Bei PI3 wurde *um die Trauben unversehrt in die Kellerei zu bringen* mit *para llevar las uvas intactas a la bodega*. Mit der wörtlichen Übersetzung des Verbs *llevar* anstatt *lleg* wird der Informationsfokus auf den Transportvorgang verlegt. Im Deutschen hingegen geht es um die Qualität des Traubenzustands. Und schließlich ist *invitar cordialmente* im spanischen Sprachgebrauch nicht üblich.

4 Schlussfolgerungen

Wie wir zu Beginn des Beitrags feststellten, ist in einer Geschäftsverhandlung die Überzeugungskraft gegenüber dem Kunden sehr wichtig, die hauptsächlich durch die richtigen Sprechakte und die Wortwahl erlangt wird. Unsere Hypothese lautete, dass KI weder zur Vorbereitung von fiktiven Verhandlungssituationen für den Dolmetschunterricht noch für den Dolmetschinsatz in Geschäftsverhandlungen geeignet ist, da ein unnatürlicher Sprachgebrauch die Verhandlungssituation und somit die kommunikative Effizienz gefährdet.

Die Ergebnisse der hier angelegten Studie bestätigen die Hypothese, dass der Sprachgebrauch bei der Vorbereitung der Verhandlungssituation im Falle der maschinellen Version die fünf für eine Geschäftsverhandlung wichtigen Dimensionen nur ungenügend erfüllt.

Es mag sein, dass eine verfeinerte Prompteingabe eine bessere Leistung von Seiten Copilots liefert, aber dies ist ein weiteres Argument dafür, dass der alleinige maschinelle Einsatz weder für die Qualität der Vorbereitung von Texten noch für das Dolmetschtraining und auch nicht für eine Verdolmetschung einer Geschäftssituation geeignet ist. Es mag auch sein, dass nicht frei zugängliche Plattformen eine bessere Leistung erbringen können. Doch abgesehen davon, scheint die hier zugrundeliegende Studie darauf hinzuweisen, dass Kommunikationstexte, bei denen es sehr auf die menschliche Interaktion ankommt, nicht ausschließlich einer KI bzw. einem Large Language Model überlassen werden sollten.

Literatur

- BRAVO MOLINA, PEDRO PABLO / BRAVO SILVA, PEDRO FERNANDO / BRAVO SILVA, DENNYS PATRICIO (2019): „La comunicación como variable determinante en el proceso de negociación.“ *Contribuciones a la Economía* 16.3. <https://www.contribucionesa laeconomia.com/index.php/contribuciones-economia/article/view/1565/1732> (8.12.2025).
- JIMÉNEZ SERRANO, ÓSCAR (2025): *La interpretación (,) A Distancia. El futuro inmediato de la profesión*. Granada: Comares.
- MOSSOP, BRIAN (2014): *Revising and Editing for Translators*. Oxon: Routledge.

- NORD, CHRISTIANE (2014): *Translating as a Purposeful Activity. Functionalist Approaches Explained*. Routledge.
- PIMENTEL SILES, MANUEL (2017): *La palabra en la ciencia de la negociación*. Tesis doctoral. Universitat Autònoma de Barcelona (España).
- REISS, KATHARINA / VERMEER, HANS J. (2014): *Towards a General Theory of Translational Action. Skopos Theory Explained*. New York: Routledge.
- VAN DER MEER, JAAP (2020): „Translation Technology – Past, Present and Future.“ *The Bloomsbury Companion to Language Industry Studies* 1: 285–310.

Paradigmenwechsel für Medientexte im Zeitalter der KI

Kopplungsanalyse von Nutzern, Themen und Affekten

Abstract

The growing influence of digital technologies is fundamentally changing communication patterns on short video platforms. Stable social communities defined by fixed affiliations and institutionalised structures are being replaced by 'light communities'. These are connected by loose forms of participation. This paper analyses User-Generated Content (UGC) comments on the Douyin and TikTok platforms. To this end, an innovative, corpus-based AI analysis framework (UTA-AI model) was developed that systematically examines platform-specific characteristics of user-generated comments along three analytical coupling levels: the user-topic, user-affect, and topic-affect relationships. The framework methodically integrates quantitative procedures (including automatic sentiment analysis using the Baidu AI API) with qualitative interpretation. The aim of this three-dimensional approach is to understand the emergence of and interaction within light communities. In addition, differences between the two platforms are identified. The results gained provide an insight into the paradigmatic change in media texts in the AI age and its influence on user behaviour. In the future, the results of this work can provide a data-driven basis for adaptive influencer marketing strategies and better targeted advertising design.

Keywords Light Communities, Nutzerverhalten, UGC-Kommentare

1 Einleitung

Mit der rasanten Verbreitung von Kurzvideos haben sich nutzergenerierte Inhalte (UGC) zu einem zentralen Medium der Alltagskommunikation entwickelt (Wang 2025: 78). Plattformen wie Douyin und TikTok schaffen durch interaktive Features, insbesondere in den Kommentarbereichen, eine neue Form partizipativer Öffentlichkeit. Hier lässt sich UGC in typische (z. B. Textkommentare) und leichtgewichtige Interaktionen (z. B. Likes) unterteilen, die distinkte „Kommentarkulturen“ oder *Light Communities* bilden. Diese werden zu einer zentralen Metrik für Engagement und sind entscheidend für plattformseitige Empfehlungen und kommerzielle Entscheidungen.

Vor diesem Hintergrund entwickelt und erprobt diese Studie ein innovatives, korpusbasiertes KI-Analyseframework (UTA-AI-Modell). Es untersucht UGC-Kommentare systematisch entlang dreier analytischer Kopplungsebenen: der Nutzer-Themen-, Nutzer-Affekt- und Themen-Affekt-Beziehung. Methodisch integriert das Framework quantitative Verfahren (u. a. automatische Sentiment-Analyse) mit qualitativer Interpretation. Das übergreifende Ziel ist es, durch den systematischen Vergleich nicht nur plattformspezifische Dynamiken zu entschlüsseln, sondern eine übertragbare wissenschaftliche Grundlage für kultursensible Plattform- und Werbegestaltung zu schaffen. Hierbei werden datengeschützte Einsichten für adaptive Influencer-Marketingstrategien gewonnen.

2 *Light Communities*

Mit dem Fortschritt der Informationstechnologie agieren Individuen im Internet nicht mehr nur als passive Rezipienten, sondern aktiv als Produzenten und VerbreiterInnen von Inhalten (Dong et al. 2023b: 40f.). Diese neuartige, partizipative Rolle bildet die Grundlage für die Entstehung spezifischer Formen von Online-Gemeinschaften, die von Blommaert mit dem Konzept der *Light Community* beschrieben wurden. Im Gegensatz zu traditionellen, stabilen sozialen Gefügen (wie Familie oder Nation) weisen *Light Communities*

keine dauerhaften, institutionalisierten Strukturen auf (Blommaert 2017: 46). Typische Praktiken – wie *Likes* oder *Retweets* – werden von Nutzenden oft als „lockerer Austausch“ wahrgenommen, können jedoch gemeinsame „emotionale Strukturen“ und so auch einen Gemeinschaftssinn fördern (ebd. 52).

Nach seiner Einführung wurde das Konzept bislang nur selten explizit als theoretische Linse in der empirischen Forschung zu Kommentarsektionen auf Kurzvideo-Plattformen genutzt. Vorhandene Studien beleuchten zwar zentrale Aspekte, die dem Konzept entsprechen:

- Massenpartizipation: Es lässt sich eine Verlagerung der diskursiven Autorität beobachten, bei der in Kommentarspalten die Stimme gewöhnlicher NutzerInnen zur dominierenden Kraft wird – ein Kernmerkmal fluiden Online-Gemeinschaften (Liu 2025: 70).
- Emotionale Ansteckung: Die Bildung kollektiver Gefühlslagen in Kommentarsektionen kann gemeinschaftsstiftende Identifikation fördern (He/Li 2025: 96).
- Gruppendynamische Abgrenzung: Der Gemeinschaftsbildung wohnt ein Moment der Abgrenzung inne, das sich in Mechanismen von Polarisierung, Ausgrenzung und Akzeptanz zeigt (Li et al. 2023: 49).

Jedoch sind diese Phänomene unter der theoretischen Perspektive der *Light Communities* bisher nur in geringem Maße wissenschaftlich erforscht worden. Genau an dieser Forschungslücke positioniert sich die vorliegende Arbeit. Ihr Ziel ist es, die Kommentarspalten von Douyin und TikTok explizit als *Light Communities* zu analysieren und die bestehenden empirischen Befunde in einen kohärenten theoretischen Rahmen zusammenzutragen. Vor diesem Hintergrund konzentriert sich diese Arbeit auf folgende Forschungsfragen:

- (1) Auf welche Weise konstituieren NutzerInnen auf den Plattformen TikTok und Douyin durch spezifische nutzergenerierte Inhalte in den Kommentarbereichen einzigartige *Light Communities*?
- (2) Welche systematischen Unterschiede lassen sich zwischen den NutzerInnen der beiden Plattformen feststellen?

- (3) Wie können die gewonnenen Erkenntnisse für inhaltsbasierte Empfehlungssysteme, Trendanalysen und kommerzielle Entscheidungsprozesse genutzt werden?

3 Korpus und KI-Analyseframework

Zur Analyse neuartiger UGC-Kommentare wurde ein Korpus mit Inhalten deutscher TiktokterInnen und Douyin-CreatorInnen aufgebaut. In den folgenden Abschnitten werden die Quellen, die Auswahlmethodik sowie die Untersuchungsdimensionen dieses Korpus detailliert erläutert.

3.1 Korpus und Datenquellen

Zur Erhebung der empirischen Daten wurde ein Korpus von UGC-Kommentaren aus den populären Kurzvideos auf den Plattformen Douyin und TikTok zusammengestellt. Die beiden Plattformen sind innerhalb bzw. außerhalb Chinas äußerst bedeutend: Bis August 2025 hatten Douyin und seine Dienste eine Gesamtzahl von 1,182 Milliarden NutzerInnen erreicht¹; und laut der Statistikplattform Statista verzeichnet TikTok weltweit 1,59 Milliarden aktive NutzerInnen, davon entfallen 21,8 Millionen auf Deutschland.²

Um eine vergleichbare und für die Forschungsfragen relevante Datengrundlage zu schaffen, wurden Kurzvideos deutschsprachiger BloggerInnen auf Douyin und TikTok ausgewählt. Die Kernkriterien sind in der folgenden Übersicht³ zusammengefasst:

Aus den 22 Profilen (12 Douyin, 10 TikTok) wurde jeweils ein populäres Video des Jahres 2025 ausgewählt. Aus deren Kommentarsektionen wurden

.....

1 Daten stammen von QuestMobile, Chinas professionellem Anbieter von Business-Intelligence-Dienstleistungen für das mobile Internet.

2 Daten stammen von „the Global Statistics“, einem spezialisierten Datenanalyseunternehmen.

3 Siehe Anhang.

alle zum Erhebungszeitpunkt in der plattformeigenen „Top“-Liste angezeigten Kommentare erfasst. Antwortkommentare wurden nur aufgenommen, sofern sie selbst in dieser Liste erschienen. Dies gewährleistet, dass das Korpus die innerhalb der Communities sichtbarsten und einflussreichsten Inhalte abbildet.

Das so konstruierte Korpus umfasst insgesamt 1045 Top-Kommentare: Das Douyin-Teilkorpus umfasst 240 Top-Kommentare; das TikTok-Teilkorpus besteht aus 805 Kommentaren.⁴

3.2 Dreidimensionale KI-Modellierung: Nutzer, Thema und Affekt

Die digitale Identität innerhalb einer *Light Community* entsteht nicht ein-dimensional, sondern in Echtzeit durch die dynamische Wechselwirkung der drei Kernkomponenten: Nutzer (U), Thema (T) und Affekt (A). Das Zusammenspiel der drei Komponenten erfasst das hinter UGC-Kommentaren liegende Verhalten anhand von Metriken wie Likes, Kommentaranzahl usw. Die Themenkomponente konzentriert sich auf die textuelle Gestalt und diskursive Relevanz der Kommentare. Daneben analysiert die Affektkomponente emotionale Tendenzen sowie visuelle Ausdrucksmittel wie Emojis.

Darauf aufbauend wird ein innovatives, multidimensionales KI-Analyseframework vorgestellt – das UTA-AI-Modell. Es integriert quantitative und qualitative Methoden durch KI-Agenten und dient als nutzerprofilbasiertes Analyseinstrument, das speziell zur Dekodierung komplexen Nutzerverhaltens auf hochinteraktiven Plattformen (z. B. Kurzvideos, soziale Medien) konzipiert ist.

.....
4 Eine detaillierte Auflistung dazu ist Anhang 1 zu entnehmen.

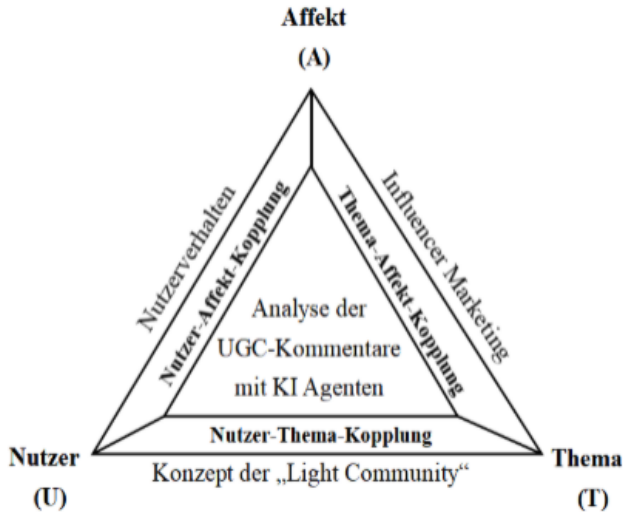


Abb. 1: UTA-AI-Modell

Die Kopplungsanalyse untersucht die Wechselwirkungen zwischen den drei Kernkomponenten:

- Die *Nutzer-Thema-Kopplung* korreliert Aufrufe, Likes und Kommentare mit Themen, um themengebundenes Langzeitengagement von situativen Interaktionen zu unterscheiden und thematische Schwerpunkte zu identifizieren.
- Die *Nutzer-Affekt-Kopplung* analysiert den Einsatz visueller Elemente, Wortspiele, Dialekte und Kommentarketten, um Interaktionsstile – kreativ-generativ, diskursiv-kontrovers und affektiv-resonant – abzuleiten.
- Die *Thema-Affekt-Kopplung* verbindet Sentimentwerte mit Kommentarlänge: Eine hohe Wortzahl bei positivem Affekt indiziert beliebte bzw. vorbildhafte Inhalte, während kurze, negative Kommentare als Frühindikatoren für Inhaltsrisiken dienen.

Das Modell ermöglicht eine integrierte Echtzeitanalyse der für die Entstehung und Dynamik digitaler Gemeinschaften entscheidenden triadischen Bezie-

hungen. Dieser Rahmen trägt zum Verständnis bei, wie digitale Partizipationskulturen entstehen und mittels datengetriebener Verfahren entschlüsselt werden können.

4 Nutzer-Thema-Kopplung

Light Communities konstituieren sich online durch spezifische Interaktionsmuster um gemeinsame Themen oder Praktiken (Dong 2023: 3). Die Kommentarsektionen von Kurzvideo-Plattformen fungieren als primärer Ort dieser thematisch fokussierten Gemeinschaftsbildung. Dabei zeigen sich die Interessenprofile der NutzerInnen auf Douyin und TikTok in einer spezifischen, oft erdichteten Textualität. Auf dieser Grundlage lassen sich thematische Schwerpunkte für ein plattformspezifisches Influencer-Marketing und Werbungsplatzierung ableiten.

4.1 Themennetzwerk

Die empirische Analyse mittels des KI-Agents *Baidu-Intelligent-Cloud* zeigt, dass NutzerInnen ausgehend vom Kernthema des Kurzvideos eigene Unterthemen generieren (siehe Anhang 1). Die darauf basierende, KI-gestützte Analyse der UGC-Kommentare visualisiert die entstehenden Diskursstränge, wie in der beigefügten Tabelle 1 dargestellt. Dabei wird deutlich, dass sich die Diskussion nicht stets auf das ursprüngliche Videothema beschränkt, sondern dynamisch auf angrenzende Themenbereiche ausweitet.

Tab. 1: Themennetzwerk

Kernthema	Unterthemen des Diskursstrangs
Fitness	Home-Workouts; Gerätenutzung; Fitness-Influencer; Follow-Along-Inhalte; Ernährung; Verletzungsrisiko
Sprachkurs	Chinesisch lernen; Aussprache; Deutsch lernen; Austausch; Sprachbarrieren; lustige Missverständnisse; Dialekte; Mehrsprachigkeit
Alltag	Ländlicher Lebensstil; Familie; Eltern-Kind-Beziehung; Alltagsgewohnheiten; kulturelle Unterschiede; Technologie im Alltag;
Kultur	Kulturvergleich Ost-West; Kultursymbole; kulturelle Identität; Kulturaustausch; kulturelle Integration; Geschichte; kollektive Erinnerung
Essen	Chinesisches Essen; ausländische Erfahrungen; Regionale Spezialitäten; Kulinarische Kulturunterschiede; emotionale Bindung
Reise	Sehenswürdigkeiten; Architektur; Interkulturelle Reiseerfahrungen; Kulinarische Entdeckungen auf Reisen; Soziales Barrieren auf Reisen
Musik	kulturelle Identität; emotionaler Ausdruck; Sprachenlernen
Mode	Mode-Influencer; Marken; Kleidungsstil; Ästhetik; Kultursymbole in der Mode

Diese Dynamik der *Light Communities* bietet ein präzises Fenster in die sich ständig neu formierenden Interessensnetzwerke der NutzerInnen. Für das Targeting in der Werbung ergibt sich daraus ein Paradigmenwechsel – von der Ausrichtung an statischen Zielgruppen hin zur kontextsensitiven Platzierung: Anzeigen können nicht nur basierend auf einem Kernthema, sondern gezielt in den sich spontan bildenden Unterthemen-Umgebungen geschaltet werden, wo die NutzerInnen bereits aktiv und emotional engagiert sind, wie z. B. das Video vom Douyin-Alltagsblogger „阿福Thomas“. Bei seinem Video prüft der Blogger das umgangssprachliche Chinesisch seines Friends durch Kommunikation mit chinesischen TCL-Angestellten innerhalb IFA2025. Während des Videos werden die aktuellen Produkte TCLs ganz und deutlich vorgestellt. Das Video dient nämlich als eine Werbung für alltägliche intelligente Möbel, was genau dem Unterthema „Technologie im Alltag“ entspricht. Dieser Ansatz ermöglicht eine deutlich höhere Relevanz und eine native Integration von Marketingbotschaften in den organischen Kommunikationsfluss der Community.

4.2 Langzeit- oder Spontanengagement

Inhalte, die an die Alltagserfahrungen der NutzerInnen anknüpfen, generieren ein signifikant höheres Interaktionsaufkommen als thematisch distanzierte Inhalte. Dies lässt sich anhand der Engagement-Daten exemplarischer Blog-

gerInnen beider Plattformen verdeutlichen (siehe Anhang 1). Die sichtbare Diskrepanz lässt sich soziokulturell erklären:

Themen mit unmittelbarer Alltagsrelevanz bieten ein breiteres Resonanzfeld für identifikatorische Bezüge und hohe Aufrufe und mehrere Likes sowie Kommentare. Beispielsweise erzielten die AlltagsbloggerInnen „Linni和Falki“ sowie „德国奶茶哥和凯蒂“ die höchsten Aufrufzahlen unter den 12 von Douyin ausgewählten Videos. Dadurch wird die Formation größerer *Light Communities* begünstigt. Dagegen stiften alltagsferne oder spezialisierte Inhalte seltener kollektiven Gesprächsstoff. Beispielsweise haben die beiden FoodlencerInnen auf TikTok mit ihren Accounts „mokka.ccino“ und „justspices“ zwar jeweils mehr als eine Million Follower, doch ihre Videos erhielten nur 28 bzw. 1 Kommentare, was auf eine bemerkenswert geringe Interaktionsrate hindeutet. Dieses Interessensprofil bedeutet, dass Lebensweltrelevanz eines Themas ein entscheidender Faktor für die Entstehung aktiver *Light Communities* ist.

Der kommerzielle Erfolg eines Influencers misst sich heute nicht allein an der Follower-Zahl, sondern an seiner Fähigkeit, fortlaufend Lebensweltrelevanz zu schaffen und daraus lebendige *Light Communities* aufzubauen. Dafür muss die Content-Strategie tief in alltägliche Praktiken und kollektive Bedürfnisse der Zielgruppe eintauchen und durch intentionales Topic-Design UGC-Interaktionen sowie Unterthemen gezielt lenken.

5 Nutzer-Affekt-Kopplung

Die NutzerInnen zeigen ihre Affekte durch den Einsatz visueller Elemente, Wortspiele und Dialekte. Dies vollzieht sich in drei charakteristischen Interaktionsmodi: einem kreativ-generativen, einem diskursiv-kontroversen und einem affektiv-resonanten Stil. Praktische Vorschläge für Influencer-Marketing und Targeting lassen sich unmittelbar aus dieser spezifischen Plattformökologie ableiten.

5.1 Affektiv-resonanter Interaktionsstil beim Einsatz visueller Elemente

Die Kommentarkultur auf Kurzvideo-Plattformen ist durch einen grundlegenden Paradigmenwechsel zum multimodalen Ausdruck gekennzeichnet, bei dem non-verbale Ressourcen wie Emojis die textbasierte Kommunikation ergänzen und intensivieren (vgl. Zhu 2022: 106).

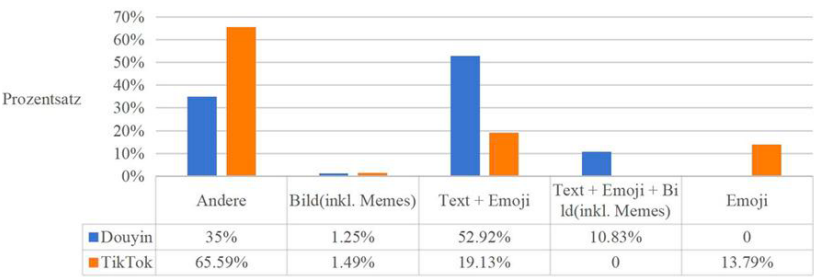


Abb. 2: Plattformspezifische Präferenz

Die korpusbasierte Analyse (Abbildung 2) zeigt eine zentrale Gemeinsamkeit: Emojis sind auf beiden Plattformen die dominierende nicht-textuelle Ressource⁵ (TikTok: 32,92 %; Douyin: 63,75 %) und dienen primär der Expressivität. Eine funktionale Konvergenz zeigt sich in der Spitzenposition: Auf TikTok dominiert das „Gesicht mit Freudentränen“ 😂 (22,90 %), auf Douyin dessen Entsprechung „Gesicht verbergen“ 🤔 (22,42 %). Beide Emojis weisen ein breites Bedeutungsspektrum auf – von Lachen über Selbstironie bis zur Resignation. Die kontextabhängige Interpretation (vgl. Anhang 3) ermöglicht dadurch schnelle affektive Resonanz.

Ein weiteres Merkmal ist der Unterschied des Einsatzes von Bildern und Memes. Auf Douyin sind sie integral (bis zu 65 % bei einigen CreatorInnen), auf TikTok dagegen selten (1,5 %). Qualitativ erfüllen sie zwei Schlüsselfunktionen:

5 Die Gesamt-Prozentzahl von Kombination „Text + Emoji“, „Text + Emoji + Bild (inkl. Memes)“ und „Emoji“.

- Deutlich sind vor allem die kontextuelle Verdichtung und Informationsergänzung (enger inhaltlicher Bezug zum Video). Sie dienen zur Veranschaulichung, wie z. B. Empfehlung konkreter Biersorten oder Pflanzensorten unter thematisch passenden Videos (vgl. Anhang 3).
- Daneben steht die Generierung geteilter visueller Räume. Die kollektiv kuratierte „Bildlandschaft“ erweitert den Originalinhalt visuell und vertieft die gruppeninterne Auseinandersetzung (Zeng/Yang 2019: 76), was den Kohäsionsgrad der *Light Community* stärkt.

Trotz ähnlicher funktionaler Grundlage (affektive Aufladung, Effizienz) unterscheiden sich die kulturelle und die semiotische Ökologie beider Plattformen grundlegend: Während TikTok einen pragmatisch-integrierten Stil pflegt, folgt Douyin einem kreativ-kurativen, hybriden Modus. Diese Differenzen spiegeln die spezifischen Technolagedesigns und Nutzerkulturen wider, welche die jeweiligen *Light Communities* maßgeblich formen.

Die Analyse verdeutlicht, dass ein affektiv-resonanter Interaktionsstil untrennbar mit der semiotischen Ökologie der Plattform verbunden ist. Daraus ergibt sich für Influencer-Marketing die wichtige Einsicht: Erfolg entsteht nicht nur durch universelle emotionale Ansprache, sondern durch plattformspezifische Interaktion und Teilhabe an den jeweiligen visuellen Symbolsystemen. Wirksame Kommunikation muss die kulturellen Stile – sei es der hybrid-kurative Modus auf Douyin oder der pragmatisch-integrierte auf TikTok – nicht nur adaptieren, sondern aktiv mitgestalten.

5.2 Kreativ-generativer Interaktionsstil durch Wortspiele und Dialekte

Die strategische Verwendung von Mehrsprachigkeit, Dialekten und syntaktischen Spielformen dient Nutzern als wesentliches Mittel zur Konstruktion sozialer Grenzen und kollektiver Identität. Dieser kreativ-generative Stil manifestiert sich in folgenden Praktiken:

- Mehrsprachigkeit als Identitätsmarker: In Kommentaren zum deutsch-türkischen Paar „sedefranz“ finden sich türkische Äußerungen (z.B. „Eşin yanında çok şükür harikasıdır“), die die kulturelle Zugehörigkeit performativ bekräftigen. Ähnlich kommentieren NutzerInnen unter dem TikTok „youneszar“ mit „Er kann Deutsch“, wodurch sie auf dessen multilinguales Medienangebot und die damit verbundene Aufmerksamkeitsökonomie anspielen.
- Dialekte als Ausdruck lokaler Verbundenheit: Kommentare unter „buonomeme“ weichen bewusst von der Standardsprache ab (z.B. „Wem ghörst du?“ – tirolerdeutsch). Durch den Einsatz solcher Dialekte inszenieren NutzerInnen eine ländliche oder regionale Identität und schaffen ein Gefühl lokaler Gemeinschaft.
- Syntaxmanipulation als exklusive „Geheimsprache“: Ein prägnantes Beispiel aus dem Douyin-Korpus ist die strategische Umstellung der Wortreihenfolge („各听我位说， 们我把顺序打乱 ...“). Diese bewusste Verletzung syntaktischer Normen erzeugt einen Code, der für MuttersprachlerInnen entschlüsselbar, für algorithmische Übersetzungstools und Nicht-MuttersprachlerInnen jedoch unzugänglich bleibt. Sie fungiert als wirksames Mittel der Abgrenzung zwischen In- und Out-Group.

Insgesamt zeigt die Analyse, dass der kreativ-generative Stil ein Akt der sozialen Grenzziehung durch kulturelle Codierung ist. NutzerInnen erschaffen durch bewusste sprachliche Normabweichungen eine exklusive semiotische Sphäre.

Vor allem bedeutet dies für Influencer-Marketing, dass wahre Relevanz dadurch entsteht, zum Architekten dieser kulturellen Codes zu werden. Ein Influencer muss die symbolischen Ressourcen seiner Community (Sprachmix, Dialekte, rhetorische Spielformen) nicht nur nutzen, sondern aktiv weiterentwickeln. Daneben folgt die Notwendigkeit einer „code-basierten Platzierung“. Anzeigen müssen diese mikro-kulturellen Sprachräume identifizieren und sich narrativ-stilistisch so einpassen, dass sie als legitimer Teil des internen Diskurses erscheinen. Erforderlich ist eine präzise Ansprache innerhalb der semiotischen Nische, nicht eine breite Zielgruppenansprache.

5.3 Diskursiv-kontroverser Interaktionsstil durch Kommentarketten

Eine besondere Form der Nutzerinteraktion in den analysierten Kommentaren zeigt sich in diskursiven und oft kontroversen Dialogketten. Diese entstehen häufig unter Beiträgen mit einem auffälligen Ungleichgewicht zwischen Like-Zahl und Antwortzahl – ein Indikator für hohe kontroverse Beteiligung statt bloßer Zustimmung. Hier wird die Like-zu-Kommentar-Ratio (L/C-Ratio) verwendet, um das Phänomen zu untersuchen.⁶

Die L/C-Ratio des derartigen Kommentars sollte die Begeisterung der Nutzerinteraktion in den Kommentaren widerspiegeln, d. h., umso niedriger diese Ratio ist, desto größer ist die Begeisterung der Nutzer für die Interaktion und kontroverser ist der Kommentar. Im Korpus finden sich aber Beispiele wie die meisten nicht aggressiven Kommentare unter „德国吃货 chihuo“ mit L/C-Ratio von 0. Unter Douyin-Blogger „德国人 Leo 乐柏说“ führte im Korpus ein Kommentar zur Architektur der Kirche mit nur 7 Likes zu 8 Antworten, die L/C-Ratio ist also 0.875 – ein Zeichen für eine fachliche oder kulturelle Debatte in einer Nische. Lange Kommentare (> 40 Zeichen), die im Korpus 7,96 % ausmachen, sind oft der Ausgangspunkt solcher Ketten. Sie enthalten typischerweise ausführliche Erzählungen, komplexe Fragen oder kritische Argumentationen, wie die lange, humorvolle Anekdote über Magenschmerzen und Instantnudeln unter „米夏德语 MiXia“ zeigt. Solche Interaktionsstile konnten einen Beitrag zu dem Influencer-Marketing und Ad-Placement leisten, darunter sind mögliche Umsetzungen der folgenden Punkte:

- Kontroverse als Engagement-Treiber: Anstatt ausschließlich auf maximale Like-Zahlen zu zielen, sollten Marken und Creator gezielt auf inhaltliche Ankerpunkte mit Diskussionspotenzial setzen. Dies können provokante Fragen, kulturelle Vergleiche oder die Darstellung ungewöhnlicher Praktiken sein (z. B. deutsche Präzision beim Gärtnern im chinesischen Kontext).

.....

6 Das Verhältnis soll aber immer größer als 0 sein; wenn es gleich 0 ist, bedeutet es, dass es unter dem Kommentar keinen einzigen Like oder eine Antwort gibt.

- Steuerung von Kommentarketten: Influencer-Marketing kann so gestaltet werden, dass es natürliche Debatten anregt. Eine Strategie ist es, mittellange bis lange Kommentare (11 bis 40+ Zeichen) von Konten zu platzieren, die eine diskursive Reaktion der Community provozieren. Das Beispiel der TCL-Diskussion unter „阿福 Thomas“ zeigt, wie Markenerwähnung organisch zu Nutzerdialogen über Produktqualität und internationale Wettbewerbsfähigkeit führen kann.
- Identifikation diskursiver *Light Communities*: Die Analyse solcher Dialogketten erlaubt es, hoch engagierte, thematisch fokussierte *Light Communities* innerhalb einer *Creator-Follower-Gemeinschaft* zu identifizieren. Für eine Werbeagentur sind diese Nutzer besonders wertvoll, da ihre aktive Beteiligung die Sichtbarkeit und Reichweite des Contents algorithmisch verstärkt.

Zusammenfassend stellt der diskursiv-kontroverse Interaktionsstil eine wirkungsvolle, aber oft unterschätzte Interaktionsform dar. Durch die strategische Generierung und Moderation von Kommentarketten können Marken authentisches Engagement fördern, Community-Bindung vertiefen und die algorithmische Verbreitung ihrer Inhalte nachhaltig unterstützen.

Die Nutzer-Affekt-Kopplung zeigt, dass der Fokus hierbei auf der Qualität und Tiefe der Interaktion liegen sollte, um ein präzises Influencer-Marketing einzuführen. Die BloggerInnen sollten auf Interaktionsstil und „code-basierte Platzierung“ achten, dadurch engere Beziehungen zu den Nutzern aufbauen und ihr eigenes kommerzielles geistiges Eigentum pflegen.

6 Thema-Affekt-Kopplung

In diesem Abschnitt werden anhand der Sentimentanalyse durch KI und der Nutzung von Interpunktion die Kommentare zu Themen untersucht, die bei den Nutzern positive oder negative Reaktionen hervorrufen. Dadurch hilft dieser Ansatz beim gezielten Marketing der BloggerInnen.

6.1 Sentimentanalyse der UGC-Kommentare

Um die Nutzerstimmung in den Kommentarbereichen beider Plattformen zu analysieren, wurden 1045 UGC-Kommentare mithilfe des KI-gestützten *Baidu-AI-API* Sentiment-Analysetools ausgewertet⁷. Die Ergebnisse sind in Tabelle 2 dargestellt. Aufgrund der komplexen Gefühlslagen und Identitäten enthalten die meisten als positiv klassifizierten Kommentare Inhalte wie kulturelle Identität, Sprachwertschätzung und Inhaltsbegeisterung. Als negativ eingestufte Kommentare beinhalten dagegen häufig humorvolle oder selbst-ironische Äußerungen, kulturvergleichende Bemerkungen sowie inhaltliche Kontroversen.

Aus dem Douyin-Teilkorpus geht hervor, dass der durchschnittliche Sentiment-Wert der aus den Douyin-Kommentaren extrahierten Texte bei 1,1667 liegt. Dies zeigt, dass die KI die insgesamt 240 Douyin-Kommentare als leicht positiv einstuft. Das TikTok-Teilkorpus zeigt eine deutliche Polarisierung der Stimmung in den Kommentarbereichen: Hier dominieren entweder überwiegend positive oder überwiegend negative Äußerungen.

Basierend auf diesem plattformübergreifenden Vergleich lässt sich für das Influencer-Marketing Folgendes ableiten: Auf Douyin sollten Marketinginhalte auf kulturelle Resonanz und emotionale Narrative setzen. Auf TikTok dagegen, mit seiner polarisierten Nutzerstimmung, ist ein differenzierter Ansatz notwendig.

.....

7 Die primären Rohdaten von den Probabilitäten und der Zuversicht werden dezimal dargestellt, während sie in Tabelle 1 als Prozentzahl gezeigt werden.

Tab. 2: Sentimentanalyse⁸

Plattform	BloggerIn	Sentiment	Zuversicht	Kommentare	Plattform	BloggerIn	Sentiment	Zuversicht	Kommentare
Douyin	Pamela Reif	2 Positiv P: 99,89%; N: 0,11%	99,76%	239	TikTok	buonomemes	0 Negativ P: 4,76%; N: 95,24%	89,41%	393
	growingannanas	0 Negativ P: 6,85%; N: 93,15%	84,77%	6458		chiaraandlaraa	2 Positiv P: 99,98%; N: 0,02%	99,95%	4
	潘咏霖 DE 讲德粤好	0 Negativ P: 17,15%; N: 82,85%	61,89%	140		easygerman	2 Positiv P: 99,99%; N: 0,01%	99,98%	6
	米夏德语 MiXia	2 Positiv P: 55,28%; N: 44,72%	0,61%	1948		Germanbull3	0 Negativ P: 14,52%; N: 85,48%	67,74%	39
	Linni 和 Falki	0 Negativ P: 2,15%; N: 97,85%	95,23%	4515		Goetheinstitut	2 Positiv P: 99,99%; N: 0,01%	99,98%	1
	德国 Moritz	2 Positiv P: 99,98%; N: 0,02%	99,96%	2601		justspices	2 Positiv P: 99,98%; N: 0,02%	99,95%	1
	德国奶茶哥和凯蒂	0 Negativ P: 1,65%; N: 98,35%	96,33%	3388		mokka.ccino	2 Positiv P: 92,66%; N: 7,34%	83,69%	28
	阿福 Thomas	2 Positiv P: 89,50%; N: 10,50%	76,66%	844		rebeccamir	2 Positiv P: 99,71%; N: 0,29%	99,35%	51
	德国人 Leo 乐柏说	2 Positiv P: 99,85%; N: 0,15%	99,67%	4650		sedefranz	2 Positiv P: 99,43%; N: 0,57%	98,32%	107
	司梵克 Volker	2 Positiv P: 99,17%; N: 0,83%	98,15%	11000		youneszar	2 Positiv P: 91,88%; N: 8,12%	81,95%	175
	德国吃货 chihuo	2 Positiv P: 92,94%; N: 7,06%	84,31%	9006					
	小梅	0 Negativ P: 1,45%; N: 98,55%	96,77%	55					

Darüber hinaus ist offensichtlich, dass die Stimmung der UGC-Kommentare einen erheblichen Einfluss auf die Verbreitung eines Videos hat. So weisen beispielsweise die beiden Videos mit den meisten Kommentaren auf Douyin eine ausgesprochen positive Stimmung auf, während die Kommentarbereiche der Videos mit den wenigsten Kommentaren negativ sind. Daher sollten BloggerInnen besonders auf das Sentiment der Kommentare achten, um die Verbreitung ihrer Videos zu erhöhen und die Ausrichtung kommerzieller Werbung zu verfeinern.

6.2 Gebrauch der Interpunktion

Die Interpunktion dient in diesem Kontext weniger der grammatischen Gliederung, sondern primär als affektiver Index. Auffällig ist die Wiederholung der expressiv kraftvollen Ausrufe- und Fragezeichen, die Emotionen verstärken.

.....

8 Der Index „Sentiment“ zeigt die Ergebnisse der Sentiment-Polaritätsklassifizierung an: 0 = negativ, 1 = neutral, 2 = positiv; die Durchschnittszahl wird um zwei Dezimalstellen gehalten. Der Index „Zuversicht“ zeigt, wie zuversichtlich die Klassifizierung des Sentiments ist.

Tab. 3: Indexikalische Zeichensetzung als Emotionsträger

Zeichensetzung	Beispielkommentar	Auswirkung
Verstärkende Einfachsetzung	擀皮的手法好厉害!	Ausruf
	YZ ODER ICH???!!!!	
Intensivierende Mehrfachsetzung	haben sie geschmeckt??	verstärkte Frage
	不是...他还没瘦回曾经, 怎么就回来了? ?	

Die Wiederholung der Zeichen konstruiert hier eine skalierbare Intensitätsstufe der Emotion, wobei mehr Zeichen größere Dringlichkeit oder Erregung kodieren.

Es lässt sich sagen, dass die emotionale Ausrichtung in den Kommentaren eine ausgeprägte Polarisierung aufweist, wobei eine stark positive Stimmung sich als vorteilhaft für die Verbreitung von Videos erweist. Daher sollten BloggerInnen darauf achten, die Stimmung in den Kommentaren zu steuern, um ein gezieltes Marketing zu ermöglichen.

7 Zusammenfassung und Ausblick

Zusammenfassend zeigt diese Arbeit, dass UGC-Kommentare unter Kurzvideos eine neuartige Diskussionsform im KI-Zeitalter der neuen Medien darstellen. Diese Form zeichnet sich durch hohe Informationsdichte, starke Emotionalität und große Interaktivität aus. Basierend auf einem multidimensionalen KI-Analyseframework untersuchte diese Studie vergleichend die Kommentarsektionen von Douyin und TikTok. Hierbei wurden anhand der drei Faktoren Nutzer, Thema und Affekt die Beziehungen und Interaktion innerhalb dieser Gemeinschaften herausgearbeitet. Darauf aufbauend lassen sich die Kernbefunde wie folgt zusammenfassen:

- Konstitution von *Light Communities*: Nutzer konstituieren *Light Communities* primär durch multimodale Interaktion (Emojis, Bilder), sprachliche Innovation (Wortspiele, Dialekte) und die Bildung affektiver Resonanzräume in thematisch fokussierten Diskursen.

- Systematische Plattformunterschiede: Es zeigt sich eine distinkte semiotische Ökologie: Douyin nutzt einen kreativ-kurativen, hybriden Modus mit externen Bildressourcen; TikTok folgt einem pragmatisch-integrierten Modus, der stärker auf plattformeigene Emojis setzt.
- Implikationen für Praxis: Die Erkenntnisse legen nahe, dass inhaltsbasierte Empfehlungssysteme und Marketingstrategien diese plattformspezifischen Interaktionsstile adaptieren müssen. Erfolg verspricht das gezielte Generieren von „Trendmomenten“, die spontane, hoch engagierte *Light Communities* auslösen und so Content-Verbreitung und Nutzerbindung steigern.

Zusammenfassend zeigt die Arbeit, dass sich Kommentarbereiche von reinen Feedback-Zonen zu dynamischen *Light Communities* entwickelt haben. Für die Entwicklung neuer Konzepte für diese sich wandelnde Medienwelt sind zwei Faktoren entscheidend. Zum einen bedarf es einer Sensibilität für die Kontexte, in welchen sich die NutzerInnen bewegen, sowie eines Verständnisses für die kulturellen Codes, welche sie in der Interaktion verwenden. Zum anderen muss ein Alltagsbezug geschaffen werden, wodurch die NutzerInnen selbst eine Verbindung mit den Inhalten aufbauen können. Hierdurch entwickelt sich die Interaktion weg von statistischen Zielgruppen hin zu einem interaktiven sich stetig weiterentwickelnden Austausch. Dieser Wandel bietet einen neuen theoretischen Zugang zum Verständnis fluiden *Light Communities* und liefert zugleich eine praktische Richtschnur für neue Konzepte zum plattformökologischen Influencer-Marketing und der Werbungsplatzierung im KI-Zeitalter.

Literatur

- BLOMMAERT, JAN (2017): „Durkheim and the Internet: on Sociolinguistics and the Sociological Imagination.“ *Tilburg Papers in Culture Studies* 173: 1–90.
- DONG, JIE (2023a): „Community Formation and Identity Construction in New Media Language: An Online Ethnography of Bilibili’s Danmu Interaction.“ *Language Policy and Planning* 1: 1–13, 177.

- DONG, JIE et al. (2023b): „Perspectives, Methods and Trends in Language and New Media Studies.“ *Chinese Journal of Language Policy and Planning* 1: 39–49.
- HE, YUTING / LI, ZI'ANG (2025): „Emotional Transmission in Cultural Tourism Government New Media and Urban Image Reconstruction: An Analysis of Sentiment Transmission Based on Comment Texts from the Guangxi Department of Culture and Tourism Douyin Account.“ *New Media Research* 7: 91–96.
- LI, JIAYI / LIU, PENG / LI, JIE (2023): „Digital Rumours and the Generation, Characteristics, and Mechanisms of Group Prejudice: An Empirical Study of Comment Texts on Douyin Videos Featuring Rural Influencers.“ *Haihe Media* 2: 44–49.
- LIU, YANIQNG (2025): „A Study on the Characteristics of Dance Commentary in Douyin Short Videos in the Self-Media Era.“ *Aesthetics* 1: 67–70.
- WANG, QI (2025): „Mainstream Media's Commitment to Values and Practical Approaches in Short Video Content Production.“ *Journalist Cradle* 9: 78–80.
- WU, MENG'EN (2023): „Analysis of the Role of User Comment Sections: A Case Study of New Media Platforms such as Douyin.“ *West China Broadcasting TV* 6: 62–64.
- YANG, YINGMEI (2022): „A Study on the Transformation of Linguistic Expression Patterns in Short Video Creation.“ *Comparative Study of Cultural Innovation* 6: 54–57.
- ZENG, GUANG / YANG, XUE (2019): „Analysis of Image Comment Usage on Sina Weibo.“ *Youth Journalist* 6: 76–77.
- ZHANG, LIGUO (2025): „Short-form Video Literary Criticism: Forging a Critical Weapon That Balances Quality and Audience Reach.“ <http://cn.chinadaily.com.cn/a/202501/21/WS678f48eca310be53ce3f2bd9.html> (05.10.2025).
- ZHANG, YAN / YANG, WENXING (2024): „A Study of Emojis in Online Short Video Comments: A Pragmatic Rhetorical Perspective.“ *Journal of Hubei Normal University (Philosophy and Social Sciences)* 6: 68–75.
- ZHU, YAOYING (2022): „Audio-Visual Integration: How Short Video Comments Combine Sound and Visual Elements.“ *News and Writing* 10: 106–110.

Anhänge

13 Davis Blogger: China und ausgewählte Videos					
Thema	BloggerIn	Followers	Veröffentlichungsdatum	Aufrufe	Zahl der Likes
Fitness	Pamela Reif	1858900	2025.09.18	1876472	15000
	growingmamas	1507000	2025.09.03	3827338	100000
Sprachkurs	潘晓芸 DE 抖音网红	29000	2025.08.18	65144	1446
	米夏德姐 MXia	53300	2025.05.23	1397188	23.000
Alltag	Lumi und Falki	1724000	2025.08.13	11201220	222000
	德国 Moritz	7812	2025.02.11	319280	9673
	德国的芬哥和闺蜜	279000	2025.09.12	9693749	163000
	阿福 Thomas	2916000	2025.09.09	426817	37000
Kultur	德国人 Leo 东柏说	3526000	2025.02.01	4111696	85900
	司梵克 Volker	665000	2025.09.21	8527943	157000
Essen	德国吃货 chhuo	113000	2025.05.31	3946100	37900
	小梅	696	2025.05.05	4656	184
10 Tik Tokler:innen und ausgewählte Videos					
Thema	Tik ToklerIn	Followers	Veröffentlichungsdatum	Aufrufe	Zahl der Likes
Alltag	youngszar	6600000	2025.07.01	289100	18600
Fitness	Germanball3	123600	2025.09.17	17300	1206
Kultur	buonommes	270000	2025.08.31	735300	80800
Reise	soderanz	109700	2025.06.14	496700	11000
Sprachkurs	enzygerman	533400	2025.06.14	456000	2517
	goetheinstitut.deutsch	280300	2025.09.03	2524	128
Essen	Modla.ecino	1600000	2025.08.23	274600	20100
	jostspees	1100000	2025.08.16	590	17
Musik	chiarandilana	61900	2025.09.02	5959	505
Mode	rebeccanir	916100	2025.07.09	546000	1730
					51

Anhang 1: Vorstellung des Korpus

Plattform	Emoji	Name	Bedeutung	Nutzungshäufigkeit
Douyin		Gesicht verbergen	Extremes Lachen, weinen vor Freude	22,42 %
		Rose	Liebe, Zuneigung, Dank, ähnlich wie [Herz]	11,21 %
		Daumen hoch	Zustimmung, Lob oder Bestätigung	6,72 %
		Gucken	Aufmerksamkeit zu schenken, aber halb im Scherz, halb im Ernst	6,28 %
		Verlegen kicherndes Gesicht	Frechheit, Schüchternheit	5,38 %
TikTok		Gesicht mit Freudentränen	Extremes Lachen, weinen vor Freude	22,90 %
		Herz	Liebe, Seele, Güte	15,49 %
		Lächelndes Gesicht mit Herzen	Zuneigung, Verehrung und manchmal Dankbarkeit	10,44 %
		Lächelndes Gesicht mit herzförmigen Augen	Gefühl, geliebt und besonders zu sein, Liebe und Dankbarkeit	3,70 %
		Daumen hoch	Zustimmung, Lob oder Bestätigung	3,70 %

Anhang 2: Emojis und Beispielkommentare

Beispiel 1	
Beispiel 2	
Beispiel 3	
Beispiel 4	“Wem hörst du”  des sag ma echt nur am Land
Beispiel 5	笑死我了，博主说的话，好像我们写的英语   [Verdammt, die Worte des BloggerInnen klingen genau wie unsere englischen Texte]
Beispiel 6	妈耶我做为中国女人都不会擀面老外居然会  [Meine Güte, selbst ich als Chinesin kann keinen Teig ausrollen, aber Ausländer können das]
Beispiel 7	
Beispiel 8	

Anhang 3: Beispielkommentare

Epilog

Thesen zur Perspektive von KI in der Fachkommunikationsforschung und -lehre

KI ist gekommen, um zu bleiben – das gilt auch für ihre Nutzung in der Fachkommunikation, deren Erforschung und Vermittlung. Der vorliegende Sammelband hat dazu projektbasiert verschiedene Facetten der Nutzung von Künstlicher Intelligenz im Rahmen der angewandten Fachkommunikation durch internationale Forschende aufgezeigt. Welche Konsequenzen lassen sich daraus ableiten und welche Tendenzen zeichnen sich für die Zukunft der Fachkommunikation mit KI ab?

Dazu sollen abschließend einige Thesen formuliert werden:

- (1) KI wird die Fachkommunikationsforschung, die Sprachvermittlung und Fachübersetzung nicht ersetzen, sondern im Rahmen der digitalen Transformation neu gestalten helfen. Alle daran Beteiligten sollten sich der KI-Dynamik nicht verschließen, sondern sie aktiv für die Bewältigung ihrer Kommunikationsaufgaben in Forschung und Lehre nutzen. Klassische Berufsbilder (ÜbersetzerIn, FachredakteurIn, SprachlehrerIn) werden ein hybrides Profil annehmen, das sprachlich-diskursive Expertise mit KI-Kompetenz, Beratung und Change-Management verbindet. Lehrkräfte nehmen die Rolle von LernarchitektInnen ein, die KI-Tutoren und automatische Feedbacksysteme für situierte fachsprachliche Kommunikation didaktisch integrieren. Personen und Institutionen, die eine derartige Transformation nicht mitvollziehen, verlieren mittelfristig an Attraktivität und Wettbewerbsfähigkeit gegenüber Nutzenden KI-basierter Mehrwertsservices in Fachkommunikation, Sprachvermittlung und Translation.
- (2) KI-Systeme sind als Hilfsmittel bzw. Werkzeuge für die Verbesserung von Fachkommunikation zu begreifen und ihre bewusste Anwendung im fachlichen Kontext kann die Kommunikation zwischen ExpertInnen, Semi-

expertInnen und Laien sowie Menschen und Algorithmen nachhaltig prägen und optimieren. KI-gestützte Tools (z. B. zur Transkription, Annotation, Korpusanalyse) werden zur Standardmethodik in Forschung und Assessment, wodurch sich der Fokus stärker auf Studiendesign, Interpretation und Theoriearbeit verschiebt. KI-basierte, adaptive Lernumgebungen mit individueller Diagnose, Feedback und Lernpfaden werden zum dominanten Paradigma in der Sprachdidaktik, insbesondere für komplexe Fertigkeiten wie Aussprache, Grammatik und Schreibtraining.

- (3) Der zielgerichtete Umgang mit KI erfordert von den AkteurInnen die Ausbildung und Anwendung von KI-Kompetenz / AI Literacy. KI-Kompetenz reift zu einer zentralen Fertigkeit, die es gilt in akademischen bzw. beruflichen Ausbildungsprogrammen zu entwickeln und in den Berufsfeldern adäquat anzuwenden. Studiengänge, die Angewandte Linguistik, Fachkommunikation, Fachübersetzen und Sprachdidaktik tangieren, müssen daher systematisch KI-Kompetenzen ausbilden, z. B. Prompting-Techniken, Evaluationsverfahren, Datenethik, Workflow-Design, und interdisziplinär mit Informatik und Recht zusammenarbeiten. Zentrale Forschungsfragen beziehen sich dabei auf die Wirksamkeit und Akzeptanz von KI in einem Fachbereich sowie die Ausgestaltung von KI-Richtlinien und didaktischen Designs für diesen Bereich.
- (4) KI wird auch in Zukunft in der Fachkommunikationsforschung und -lehre vor allem in drei wichtigen Funktionen wirksam:
 - (a) als datengestütztes Analyse- und Produktionsinstrument für große mündliche und schriftliche Kommunikate, um deren sprachliche Dynamik zu beobachten, zu analysieren und zu optimieren;
 - (b) als Tool, indem es Kommunikationsprozesse vorbereitet, begleitet und evaluiert, z. B. die Automatisierung von Textrezeptions- und -produktionsprozessen durch Large-Language-Model-basierte Workflows oder die Personalisierung der Informationsverarbeitung;
 - (c) als Tutor zur Anleitung, Begleitung und Bewertung von Kommunikationsprozessen in Bezug auf Wirksamkeit und Effizienz, z. B. in der studentischen Ausbildung im Selbststudium, in Form von Learning Analytics, der Ausweitung adaptiver Schreib-, Tutoring- und Feedback-Systeme. Sprachlernen bleibt aber trotz KI-Agenten und Echt-

zeitübersetzung relevant, verschiebt den Fokus jedoch stärker auf interkulturelle, diskursive und berufsbezogene Handlungsfähigkeit.

- (5) Kuratorische, didaktische und diskursgestaltende Rollen von KI werden in Zukunft an Bedeutung gewinnen und die Kommunikation begleiten. Zu Fairness, Bias, Datensicherheit, Urheberrecht und der Verschiebung von Autorschafts- und Bewertungspraktiken sind Bewertungsmaßstäbe und Regelungen zu treffen. In diesem Kontext werden neue, interdisziplinäre Forschungsfelder entstehen. Die Rolle der Angewandten Linguistik und damit der Fachkommunikationsforschung wird sich unter der Dynamik von KI stark verändern, indem sie zu einem zentralen Akteur in der Aushandlung von Sprach- und Wissenspolitik im KI-Zeitalter reift.

Über die Autorinnen und Autoren / About the Authors

Beneharo Álvarez-Pérez (Universidad de Las Palmas de Gran Canaria) is a full-time lecturer at the University of Las Palmas de Gran Canaria (ULPGC), where he is a member of the TrIAl (IDeTIC) research group. His research focuses on interpreter pre-task preparation processes, with particular attention to terminology management and the integration of Artificial Intelligence into interpreter education.

ORCID: 0000-0002-4916-882X

Gennaro Ascione (Université Bourgogne Europe) is a PhD candidate at the Université Bourgogne Europe and the University of Naples “Parthenope”. His research focuses on the linguistic and discursive evolution of the global wine industry through Organisation Internationale de la Vigne et du Vin (OIV). His interests include terminology, specialised languages, lexical standardisation, diachronic methodologies, and digital humanities.

ORCID: 0009-0007-6994-3460

Sandra Braun (Westfälische Hochschule Zwickau) ist Lehrkraft für besondere Aufgaben in Deutsch als Fremd- und Fachsprache und wissenschaftliche Mitarbeiterin im DAAD-Projekt E-Candee an der WHZ (Zwickau). Sie studierte Theater-, Film- und Medienwissenschaft an der Universität Wien und absolvierte dort eine DaF/DaZ-Ausbildung. Internationale Lehr- und Projekterfahrungen führten sie u. a. nach Aserbaidschan, Kasachstan, Tschechien (OeAD-Lektorin an der TU Liberec 2017–2022) und Österreich. Forschungsschwerpunkte: Fachsprache, Sprachenpolitik, Plurizentrik, Methodik und Didaktik.

ORCID: 0000-0002-1086-7041

Ines-A. Busch-Lauer (Westfälische Hochschule Zwickau) studierte Anglistik, Russistik, Psychologie und Pädagogik an der Universität Leipzig (1978–1983), promovierte in Kontrastiver Fachtextlinguistik (1986) und habilitierte sich im Jahr 2000. Seit 2006 ist sie Professorin für Fachenglisch/Kommunikation an der WHZ (Zwickau) und leitet den Bereich Fachbezogene Sprachausbildung. Forschungsschwerpunkte: Fachkommunikation, Didaktik der Fachsprachen, Künstliche Intelligenz. Sie ist Herausgeberin der Reihe *Fach, Sprache und Kultur* im Verlag Frank & Timme.

ORCID: 0000-0002-8087-911X

José Tomás Conde Ruano (University of the Basque Country, UPV/EHU) holds a PhD in Translation and Interpreting from the University of Granada and is currently Senior Lecturer at the University of the Basque Country. His research focuses on translators' professional practices, translation quality and assessment, grounded in cognitive translology and empirical methods. He has collaborated with leading Spanish research groups in the field, including PACTE, PETRA, GENTT, TRALIMA/ITZULIK, and ECIS+.

ORCID: 0000-0002-4089-5019

Christophe Cruz (Université Bourgogne Europe) is a Full Professor of Computer Science at the Université Bourgogne Europe and Head of the Computer Science Department at IUT Dijon-Auxerre. His research spans neuro-symbolic AI, knowledge graphs, and agentic AI: developing autonomous, decision-making multi-agent systems that combine symbolic reasoning, planning and adaptive learning.

ORCID: 0000-0002-5611-9479

Isabella Ferron (Università degli Studi di Modena e Reggio Emilia) studierte an der Ca' Foscari Universität Venedig sowie der Eberhard-Karls-Universität Tübingen Germanistik, Anglistik und Sprachwissenschaft. Ihre Promotion erfolgte im Cotutelleverfahren an der LMU und der Universität Padua. Sie setzte ihre akademische Laufbahn mit Forschungsaufenthalten in Potsdam, Berlin und London fort. Sie war Lehrbeauftragte in Venedig, Padua, Siena, Mailand und Leipzig und wissenschaftliche Mitarbeiterin in Rom. Derzeit lehrt sie

Deutsche Sprache, Sprach- und Übersetzungswissenschaft an der Universität Modena und Reggio Emilia. Forschungsschwerpunkte: kognitive und kontrastive Linguistik, historische Semantik, Sprachgeschichte, Politolinguistik.

ORCID: 0000-0003-2327-7603

Iryna Gaman (Nationale Technische Universität der Ukraine „Ihor Sikorskyj Kyjiwer Polytechnisches Institut“) studierte an der Fakultät für Linguistik der Nationalen Technischen Universität der Ukraine „KPI“ und ist seit 2009 Dozentin (Übersetzerin/ Dolmetscherin und Hochschullehrerin) für deutsche und englische Sprache am Lehrstuhl für Theorie, Praxis und Übersetzung der deutschen Sprache an der Fakultät für Linguistik der KPI. Masterstudium an der Europa-Universität Viadrina Frankfurt (Oder) im Studiengang „Intercultural Communication Studies“ (2012); Promotion zu „Linguostilistische Charakteristika der Texte deutschsprachiger Zooführer“, Fachgebiet „Germanische Sprachen“ (Kyjiw, 2020). Forschungsinteressen: Stilistik, Textsortenlinguistik, interkulturelle Kommunikation, konsekutives Dolmetschen, Phraseologie.

ORCID: 0000-0002-9042-5162

Laurent Gautier (Université Bourgogne Europe) is a Full Professor for German and Applied Linguistics at the Université Bourgogne Europe, where he has headed the TIL research laboratory (2018–2026). His work focuses on linguistic approaches to specialised discourse, particularly professional discourse. From a methodological point of view, he is working at the intersection of cognition and computer modelling. The fields explored to date cover agri-food, economics, law, and meteorology.

ORCID: 0000-0002-6210-410X

Carlos J. Guerrero Ramos (Universidad de Granada) hat in Granada, Cambridge und Heidelberg Übersetzen und Dolmetschen, Spanische Philologie und Anglistik studiert. Nach früherer Lehrtätigkeit an der Thames Valley University in London, an der Universität zu Köln, an der Technischen Hochschule Köln und an der Universität Pablo de Olavide in Sevilla ist er seit 2008 an der Fakultät für Übersetzen und Dolmetschen der Universität Granada beschäftigt. Hier unterrichtet er theoretische und praktische Kurse des Über-

setzens und Dolmetschens für die Sprachkombinationen Deutsch–Spanisch und Englisch–Spanisch.

ORCID: 0000-0003-3196-5326

Yvonne Hörmann, Dr. M.A., (MB-Dienststelle für die Realschulen in Mittelfranken) ist Realschullehrerin für Deutsch, Musik und IT sowie Sachbearbeiterin beim Ministerialbeauftragten in Mittelfranken. Seit über 15 Jahren engagiert sie sich im Arbeitskreis Leseförderung (AK #lesen.bayern) am Staatsinstitut für Schulqualität und Bildungsforschung (ISB). Zudem ist sie als Referentin in der Lehrerfortbildung tätig und seit 2008 mit Lehraufträgen an Hochschulen vertreten. Themen- und Forschungsschwerpunkte: Lesemethoden, Umgang mit Textstrukturen, Concept Maps, innovativer Einsatz von KI zur Lese- und Sprachförderung.

ORCID: 0009-0002-8288-8172

Emilia Iglesias Fernández (Universidad de Granada) is Senior Lecturer in Interpreting at the University of Granada. She holds a PhD in Interpreting, has been involved in projects related to the role of paralinguistic cues in the interpreter's vocal behaviour in conference interpreting. She has also explored the import of paralinguistic features in the audio descriptor's voice in audio description of films for the visually impaired and in Telephone Interpreting. Participated in the EU Project *SHaping the interpreters of today and of the future: SHIFT in Orality*. She has analysed the import of the prosodic features displayed by telephone interpreters in their interactions with service providers and with users using a TI corpus. Current research: role of AI instruments in interpreting (e.g. Copilot) both in training and professional services.

ORCID: 0000-0002-4971-6347

Guldastan Ismailova (Kirgisische Staatliche Technische Universität I. Razzakov (KSTU), Bischkek) ist Deutschdozentin am Kirgisisch-Deutschen Technischen Institut (DKTI) an der Kirgisischen Staatlichen I. Razzakov Universität, Bischkek, insbesondere für „Technisches Deutsch in der Informatik“. Sie erarbeitet im Tandem mit der WHZ (Zwickau) ein Blended-Learning-Programm zu Fachdeutsch Informatik, in dem bereits vier Lehrbücher in Zusammenar-

beit mit der DAAD-Lektorin Anja Lange entstanden und mit Genehmigung des Ministeriums für Bildung und Wissenschaft der Kirgisischen Republik veröffentlicht wurden. Arbeitsschwerpunkte: Weiterbildung für DeutschlehrerInnen in Bischkek, am Goethe-Institut in Almaty (Kasachstan), am Tashkent-Institut für Architektur und Universität Samarkand (Usbekistan).

ORCID: 0000-0002-8592-5517

Filip Kalaš (Wirtschaftsuniversität in Bratislava, Fakultät für Angewandte Sprachen) ist Linguist und beschäftigt sich in seiner Forschung mit Fachphraseologie, Kollokationen und Metaphern im Finanz- und Wirtschaftsdiskurs. Er kombiniert Korpuslinguistik mit angewandten statistischen Methoden, um lexikalische Muster und konzeptuelle Strukturen zu untersuchen. In seinen aktuellen Projekten erforscht er das Potenzial und die Grenzen großer Sprachmodelle in der Sprachforschung, wobei er sich auf KI-basierte Textanalyse und die Übersetzung von Redewendungen in verschiedenen Fachregistern konzentriert.

ORCID: 0000-0003-2554-9465

Anja Anna Käuper (Universidad de Granada) ist Übersetzerin und Dolmetscherin für die Sprachkombination Deutsch–Spanisch und Dozentin an der Fakultät für Übersetzen und Dolmetschen der Universität Granada. Zuvor hat sie bereits an der Universität zu Köln und der Universität Pablo de Olavide in Sevilla gelehrt und geforscht, wobei die Schwerpunkte ihrer Arbeit immer die hispanistische Sprachwissenschaft und Übersetzen und Dolmetschen waren.

ORCID: 0000-0002-3571-2003

Christian Lang (Leibniz-Institut für Deutsche Sprache, Mannheim) is a linguist and works in the Department of Grammar at the Leibniz Institute for the German Language. He was involved in the development and restructuring of the department's terminological resources and has recently focused primarily on researching large language models. In particular, he is interested in the use of LLMs for automatic question answering and in how the factual accuracy of generated responses can be optimised.

ORCID: 0009-0003-2916-4658

Anja Lange (Westsächsische Hochschule Zwickau) studierte West- und Ostslawistik an der Universität Leipzig und an der Nationalen Taras-Ševčenko-Universität Kyjiw. Sie war 2013/2014 DAAD-Sprachassistentin an der Nationalen Kyjiw-Mohyla-Akademie und zwischen 2014 und 2019 DAAD-Lektorin an der Nationalen Technischen Universität der Ukraine „Kyjiwer Polytechnisches Institut benannt nach Ihor Sikorsky“. Zwischen 2020 und 2025 war sie DAAD-Lektorin in Bischkek, Kirgistan. Seit 2025 ist sie Dozentin für Deutsch als Fremdsprache an der WHZ (Zwickau). Forschungsschwerpunkte: Gender in der ukrainischen Kultur- und Literaturgeschichte, in der ukrainischen Identitätsbildung, Landeskunde- und Fachsprachenforschung DaF.

ORCID: 0000-0001-5218-7485

Jie Li (Nankai University, China) ist Philologin und promovierte an der Universität Leipzig. Derzeit arbeitet sie als Assoziierte Professorin an der Fakultät für Fremdsprachen der Nankai University in Tianjin, China. Forschungsschwerpunkte: deutsch-chinesische Übersetzung, kognitive Translatologie, interkulturelle Studien und Praxis.

ORCID: 0000-0001-9740-7808

Peiyu Liang (North Western Polytechnical University, China) ist Masterstudentin an der Northwestern Polytechnical University in Xi'an und forscht im Bereich der germanistischen Translatologie.

Dejun Liu (North Western Polytechnical University, China) ist Masterstudent an der Northwestern Polytechnical University in Xi'an und forscht im Bereich der germanistischen Translatologie.

Pavol Lipták (Wirtschaftsuniversität in Bratislava, Handelsfakultät) konzentriert sich auf die Schnittstelle zwischen Category-Management, Vertriebsstrategien, Marketing und Wirtschaftslinguistik. Seine aktuelle Forschung untersucht, wie Sprache als Verkaufsinstrument fungiert und wie KI-gestützte Analysen im Einzelhandel Kommunikation in messbare Geschäftsergebnisse umwandeln. Er arbeitet an Projekten mit, die sich mit mehrsprachigem

Marketing, digitaler Überzeugungskraft und der Rolle sprachlicher Signale in Online-Preisstrategien befassen.

ORCID: 0009-0009-1802-1700

Rui Ma (Universidad de Granada) is Professor of Chinese at the Faculty of Translation and Interpreting, University of Granada. She holds a PhD with international mention specialising in Spanish–Chinese translation and interpreting, intercultural communication, tourism-related localisation. Her research also explores innovative pedagogies, transcreation and terminology. In 2023, she contributed to projects supported by the Chinese Ministry of Education and the Sichuan Provincial Federation of Social Sciences.

ORCID: 0000-0002-5360-7581

Margret Mundorf ist selbstständige Linguistin, Trainerin und Dozentin mit den Themenschwerpunkten Schreiben in Wissenschaft und Beruf im Kontext von KI sowie Fach- und Wissenskommunikation in Bildung, Wirtschaft und Recht. Sie ist Mitgründerin des „Virtuellen Kompetenzzentrums: Künstliche Intelligenz in Bildung, Wissenschaft & Arbeitswelt“ (VK:KIWA), Vorstandsmitglied der Gesellschaft für Schreibdidaktik und Schreibforschung (gefsus) und Mitglied des Heidelberger Arbeitskreises für Rechtslinguistik. Sie promoviert an der Schnittstelle von Rechtslinguistik und KI-Forschung an der Universität Heidelberg.

ORCID: 0000-0001-7002-6233

Waldemar Nazarov (Université Bourgogne Europe, Centre Interlangues – Texte, Image, Langage (TIL)/Johannes Gutenberg-Universität Mainz, Romanisches Seminar) ist deutsch-französischer Brückenprofessor am Romanischen Seminar der Johannes Gutenberg-Universität Mainz und am Centre Interlangues TIL der Université Bourgogne Europe. Sein binationaler Lehrstuhl (*Chaire franco-allemande*) wird von der Deutsch-Französischen Hochschule gefördert und befasst sich mit der Angewandten Linguistik, Translationswissenschaft und Interkulturalität. Die Forschungsschwerpunkte liegen in der Fachsprachenforschung, mit besonderem Fokus auf der Rechtslinguistik,

sowie der kulturwissenschaftlich orientierten Sprachwissenschaft und Translatologie.

ORCID: 0009-0002-3674-2551

Jessica Pérez-Luzardo Díaz (Universidad de Las Palmas de Gran Canaria) is Associate Professor, accredited as Full Professor, at the University of Las Palmas de Gran Canaria (ULPGC), where she is a member of the TrIAL (IDeTIC) and ECIS+ (University of Granada) research groups. Her work focuses on interpreter training, simultaneous interpreting pedagogy, cognitive load, and affective factors such as anxiety, motivation, and self-concept. She also explores the integration of Artificial Intelligence and virtual reality in interpreter education.

ORCID: 0000-0002-2748-9892

E. Macarena Pradas Macías (Universidad de Granada) is Senior Lecturer at the University of Granada and accredited Professor. She began her research career in 1999, shortly after starting her teaching at the University of Granada. PhD on the quality parameter and fluency (2003). Research activities: identification of the pause pattern in simultaneous interpreting; innovative proposals in the field of interpreting quality (influence of non-verbal quality parameters) and gender in text production and its perception. Since March 2023, Researcher-in-Charge of the ECIS+ research group (Evaluation of Quality in Simultaneous Interpreting and Related Fields) at the University of Granada.

ORCID: 0000-0002-1834-4907

Thorsten Roelcke (Technische Universität Berlin) ist Professor für Deutsch als Fremdsprache mit den Schwerpunkten Linguistik und Didaktik von Fachsprachen und Interkulturelle Kommunikation an der TU Berlin. Seine Forschungsschwerpunkte liegen in den Bereichen Fachsprachen und berufliche Kommunikation, Deutsch als Fremdsprache und Mehrsprachigkeit, Sprachwandel und Sprachgeschichte des Deutschen.

ORCID: 0000-0002-3894-6075

Elisa Serna-Martínez (Universidad de Granada) is Assistant Professor in Translation and Interpreting at the University of Granada. As member of the ECIS+ research group on quality in interpreting and translation, her work focuses on interpreter training, embodied approaches to listening, and the impact of AI tools on interpreting and literary translation. She also works on collaborative translation and decolonial, gender-oriented translation practices, combining classroom innovation with practice-based research on cognition, language perception and professional ethics.

ORCID: 0000-0001-9821-1049

Karolina Suchowolec (Technische Hochschule Köln) is a linguist and Professor of Terminology and Multilingual Specialized Communication at TH Köln. Her research and teaching focus on various aspects of terminology and terminology management across multiple domains, covering conceptual, technological, and procedural dimensions.

ORCID: 0009-0007-6321-219X

Sarah Theroine (Université Bourgogne Europe) is a researcher in linguistics and NLP whose work focuses on corpus construction, controlled languages, and domain-specific text processing. She completed a CIFRE PhD at the Université Bourgogne Europe (TIL/ICB) in partnership with DAVI The Humanizers, developing methods and tools for assessing idiomaticity and comparability in specialised comparable corpora. Her research includes LLMs optimisation, terminology studies, prompt-engineering and frame semantics.

ORCID: 0000-0003-4238-763X

Janina M. Vernal Schmidt (Westsächsische Hochschule Zwickau) arbeitet an der Fakultät Angewandte Sprachen und Interkulturelle Kommunikation als Professorin für Deutsch als Fremd- und Fachsprache (Wirtschaftskommunikation Deutsch im multilingualen Kontext). Forschungsschwerpunkte: kritisch-engagierte Wirtschaftskommunikation Deutsch, hochschulische Internationalisierung, Diversity und Bildungsmedien im Fremdsprachenunterricht, Hochschulpolitik und rekonstruktive Forschung.

ORCID: 0000-0003-4832-5907

Tanja Wissik (Österreichische Akademie der Wissenschaften) ist Senior Scientist an der Österreichischen Akademie der Wissenschaften und unterrichtet an den Universitäten Wien und Graz. Sie studierte Übersetzen und Dolmetschen an der Universität Graz und promovierte an der Universität Wien im Fach Translationswissenschaft mit besonderer Berücksichtigung der Terminologiewissenschaft und Korpuslinguistik. Forschungsschwerpunkte: Terminologiewissenschaft, digitale Sprachressourcen und Sprachtechnologien, insbesondere Terminologische Ressourcen und Korpora, Veränderung von Terminologieworkflows durch Technologien.

ORCID: 0000-0002-1631-4560