

MASTERARBEIT | MASTER'S THESIS

Titel | Title

Der Einfluss von Unsicherheitsmarkern und digitaler Reife auf
die Fehlererkennung in KI-Antworten

verfasst von | submitted by

Pauline Kristina Marie Schmidt BSc

angestrebter akademischer Grad | in partial fulfilment of the requirements for the degree of
Master of Science (MSc)

Wien | Vienna, 2025

Studienkennzahl lt. Studienblatt | Degree
programme code as it appears on the
student record sheet:

UA 066 840

Studienrichtung lt. Studienblatt | Degree
programme as it appears on the student
record sheet:

Masterstudium Psychologie

Betreut von | Supervisor:

Univ.-Prof. Dipl.-Psych. Dr. Arnd Florack

Inhaltsverzeichnis

Abstract	2
Der Einfluss von Unsicherheitsmarkern und digitaler Reife auf die Fehlererkennung in KI-Antworten	3
Theoretischer Hintergrund	6
<i>Informationsbezug durch digitale Medien</i>	<i>6</i>
<i>Das Zeitalter der Chatbots</i>	<i>8</i>
<i>Psychologische Einflussfaktoren im Interaktionsdesign von Chatbots.....</i>	<i>10</i>
<i>Unsicherheitsmarker und ihre Wirkung.....</i>	<i>13</i>
<i>Die Rolle der Digitalen Reife im Kontext des der Chatbot Nutzung.....</i>	<i>14</i>
<i>Die Studie.....</i>	<i>16</i>
<i>Operationalisierung der Variablen.....</i>	<i>17</i>
Methodik	18
<i>Stichprobenumfang und Power-Analyse.....</i>	<i>18</i>
<i>Teilnehmende.....</i>	<i>18</i>
<i>Design.....</i>	<i>19</i>
<i>Prozedur.....</i>	<i>20</i>
<i>Statistische Analysen</i>	<i>22</i>
Ergebnisse.....	22
<i>Korrelationen</i>	<i>22</i>
<i>Linear Mixed Model</i>	<i>24</i>
Diskussion.....	26
<i>Praktische Implikationen</i>	<i>27</i>
<i>Limitationen</i>	<i>28</i>
<i>Ausblick</i>	<i>29</i>
Literatur	31
Anhang	38

Abstract

Im Kontext der zunehmenden Nutzung von auf künstlicher Intelligenz (KI) basierten Sprachmodellen stellt sich die Frage, wie Nutzer*innen mit potenziell fehlerhaften Chatbot-Antworten umgehen. Ziel der vorliegenden Arbeit war es, den Einfluss von sprachlichen Unsicherheitsmarkern sowie individueller *digitaler Reife* auf Vertrauen und den Umgang mit Chatbot-Antworten zu untersuchen. Insgesamt nahmen 69 Studierende an der Laborstudie teil. Sie erhielten 20 Screenshots von Chatbot-Dialogen, die sich auf ein wissenschaftliches Paper bezogen – jeweils 10 Antworten mit Konjunktivformulierungen (Unsicherheitsmarker) und 10 mit einem allgemeinen Warnhinweis (Kontrollbedingung), wie er im weitverbreiteten Chatbot ChatGPT vorkommt. Jeweils die Hälfte der Antworten enthielten gezielt eingebettete Falschinformationen. Die Teilnehmenden beurteilten die Vertrauenswürdigkeit der Antworten, entschieden, ob sie diese ohne eigene Recherche akzeptierten oder im Original-Paper nachlasen, und bearbeiteten anschließend Fragen zur Vertrauenswürdigkeit der Chatbots und zur eigenen *digitalen Reife*. Die Ergebnisse zeigen, dass Konjunktivformulierungen im Vergleich zu Warnhinweisen zu signifikant kritischeren Bewertungen der Chatbot-Antworten führten. Auch das Vertrauen in die Chatbots war in der Konjunktiv-Bedingung geringer. Das gewählte Vorgehen zur Fehlerprüfung (Nachlesen im Text oder direkte Übernahme von Chatbot-Antworten) wurde hingegen nicht signifikant beeinflusst. Es zeigte sich weder im linearen gemischten Modell noch in den bivariaten Korrelationen ein signifikanter Zusammenhang zwischen *digitaler Reife* und der Bewertung der Chatbots. Die Befunde der vorliegenden Arbeit liefern praktische Implikationen für das Design vertrauenswürdiger KI-Systeme und verdeutlichen, wie die sprachliche Gestaltung sowie interindividuelle Unterschiede das Nutzerverhalten beeinflussen können.

Keywords: Interaktionsdesign, künstliche Intelligenz, Digitale Reife, KI-Halluzinationen, Fake News, Trust in Automation

Der Einfluss von Unsicherheitsmarkern und digitaler Reife auf die Fehlererkennung in KI-Antworten

Künstliche Intelligenz (KI) in Form von Chatbots wie ChatGPT erfreut sich weltweit wachsender Beliebtheit, insbesondere wenn es um das schnelle Beziehen von Informationen geht. Eine Marktforschungsstudie von Forbes Advisor und OnePoll (Haan, 2023) mit 2.000 US-amerikanischen Teilnehmer*innen ergab, dass 67 % der Befragten KI-Tools wie ChatGPT traditionellen Suchmaschinen vorziehen. Dies verdeutlicht, wie schnell sich KI-Chatbots wie ChatGPT, welcher seit Ende 2022 frei verfügbar ist, zur Informationsbeschaffung etabliert haben.

Diese Tendenz birgt jedoch Risiken, da KI-Systeme fehlerhafte und sogar erfundene Informationen generieren können (Bhattacharyya et al., 2023; Borji, 2023; van Swol, Florack, & Reimer, 2025; Yoo, Kang, & Oh, 2024). In einer Studie von Bhattacharyya et al. (2023) wurde die inhaltliche Korrektheit von Referenzen in medizinischen Fachtexten, die von ChatGPT-3.5 generiert wurden, untersucht. Auf standardisierte Eingabeaufforderungen, auch *Prompts* genannt, hin ließen sie 30 kurze Fachartikel mit jeweils mindestens drei Literaturangaben generieren. Hierbei waren lediglich 7 % der Ergebnisse stichhaltig, von den übrigen zitierten Quellen waren 47 % frei erfunden und 46 % existierende Quellen, die inhaltlich allerdings falsch wiedergegeben wurden.

Bei Chatbots wie ChatGPT von OpenAI handelt es sich um sogenannte Large Language Models (LLM), die mit großen Mengen an Textdaten trainiert werden, um sprachliche Muster zu lernen. Sie kennen keine spezifischen Fakten, sondern basieren auf Wahrscheinlichkeiten und ahmen menschenähnliche Konversationen nach, dadurch ergibt sich ein hohes Risiko, dass plausibel klingende, aber fehlerhafte Inhalte entstehen, die unkritisch übernommen werden (Borji, 2023).

Stanley, Whitehead und Marsh (2022) sowie Kiziler (2024) warnen davor, dass die klar verständliche Darstellung der Inhalte dazu führen kann, dass sie von Nutzer*innen ungeprüft übernommen und wie Fakten behandelt werden. Zusätzlich legen auch klassische Modelle der Informationsverarbeitung nahe, dass verständlich präsentierte Inhalte mit höherer Wahrscheinlichkeit akzeptiert werden – unabhängig davon, ob sie korrekt sind.

Eine psychologische Erklärung für dieses Phänomen liefern Baruch Spinoza (1677/1982) und Dan Gilbert (1991). Spinoza nahm an, dass der Verstehensprozess automatisch mit einem Glaubensakt einhergeht: Menschen glauben zunächst das, was sie verstehen – wohingegen das Ablehnen einer Information zusätzliche kognitive Anstrengung erfordert. Gilbert bestätigte diese Annahme empirisch und konnte zeigen, dass Menschen Aussagen erst dann hinterfragen, wenn sie bewusst Widersprüche wahrnehmen oder gezielt zum Nachdenken angeregt werden. Da Chatbot-Dialoge in ihrer Struktur und Formulierung der menschlichen Kommunikation stark ähneln, besteht die Möglichkeit, dass deren Aussagen allein aufgrund ihrer sprachlichen Klarheit und inhaltlicher Konsistenz mit zuvor korrekten Informationen als glaubwürdig wahrgenommen werden. Eine kritische Prüfung – etwa durch Quellenvergleiche oder ergänzende Recherchen – bleibt dabei häufig aus. Erst wenn Nutzer*innen aktiv zur Reflexion angeregt werden oder auf Widersprüche stoßen, erfolgt eine systematische Auseinandersetzung mit dem Inhalt und gegebenenfalls eine Korrektur der ursprünglichen Einschätzung.

Welche Personenfaktoren dazu beitragen, dass Nutzer*innen kritisch mit KI-generierten Inhalten umgehen, statt diese passiv zu nutzen, ist bislang kaum empirisch untersucht worden. Eine mögliche Einflussgröße könnte die digitale Reife sein, systematische Befunde hierzu stehen jedoch weitgehend aus (van Swol et al., 2025). Digitale Reife beschreibt die Fähigkeit, digitale Technologien reflektiert und kompetent zu nutzen, Inhalte kritisch zu hinterfragen und informierte Entscheidungen im digitalen Kontext zu treffen

(Laaber et al., 2023). Personen mit höherer digitaler Reife könnten Chatbots aktiver und kritischer verwenden (van Swol et al., 2025) und somit auch potenzielle Fehler in KI-Antworten häufiger erkennen. Ziel dieser Arbeit ist es, einen Beitrag zu der Beantwortung der Frage zu leisten, wie ein verantwortungsvoller Umgang mit KI und potenziell fehlerhaften Inhalten durch gezieltes Interaktionsdesign gefördert werden kann.

Interaktionsdesign beschreibt die Gestaltung der Interaktion zwischen Nutzer*in und System und unterstützt die zielgerichtete Nutzung digitaler Anwendungen (Teo, 2024). Ein Beispiel für ein solches Gestaltungsmerkmal, das zu einem verantwortungsvolleren Umgang führen soll, ist der Warnhinweis bei ChatGPT unter dem Eingabefeld, der lautet: „ChatGPT kann Fehler machen. Überprüfe wichtige Informationen“ (OpenAI, 2025).

Da die Antworten von LLMs auf Wahrscheinlichkeiten basieren, liegt der Gedanke nahe, dass ein Unsicherheitsmarker bei Antworten mit einer unterdurchschnittlichen Richtigkeitswahrscheinlichkeit eingesetzt werden könnte. Dadurch würde der Chatbot, ähnlich wie Menschen es in der Kommunikation tun, seine Ungewissheit mitteilen und einen Faktencheck anregen. Yoo et al. (2024) konnten erstmals in einer Studie beobachten, dass rhetorische Mittel eine entscheidende Rolle für die Glaubwürdigkeit von KI-Antworten spielen könnten. Sie beobachteten, dass kurze technische Antworten eher als richtig wahrgenommen wurden als solche mit ausgeprägten rhetorischen Mitteln. Dies unterstützt die Annahme, dass sprachliche Unsicherheitsmarker einen wichtigen Beitrag zur Fehlererkennung leisten könnten.

Aus diesem Kontext ergibt sich folgende Forschungsfrage: *Wie lassen sich potenziell falsche Inhalte durch sprachliche oder kontextuelle Hinweise besser als solche erkennbar machen, um Fehlinformationen entgegenzuwirken und einen verantwortungsvollen Umgang mit Chatbots zu fördern?*

Diese Forschungsfrage basiert auf der Annahme, dass sowohl sprachliche Gestaltungselemente wie Unsicherheitsmarker als auch individuelle Unterschiede, insbesondere im Hinblick auf digitale Reife, beeinflussen, wie Nutzer*innen mit fehlerhaften KI-Inhalten umgehen. Frühere Modelle der Informationsverarbeitung (z. B. Petty & Cacioppo, 1986; Strack & Deutsch, 2004) zeigen, dass verständlich formulierte Informationen häufig ungeprüft übernommen werden – selbst dann, wenn sie inhaltlich nicht korrekt sind. Gleichzeitig legen aktuelle Befunde nahe, dass Personen mit höherer *digitaler Reife* KI-generierte Aussagen kritischer hinterfragen (Laaber et al., 2023; Zimmermann, 2025). Daraus ergibt sich ein bislang wenig erforschter, aber gesellschaftlich relevanter Aspekt: die Frage, ob durch gezielte sprachliche Hinweise und individuelle Kompetenzen ein bewusster und reflektierter Umgang mit Chatbot-Antworten gefördert werden kann.

Theoretischer Hintergrund

Informationsbezug durch digitale Medien

Die technologischen Fortschritte im Zuge der Digitalisierung haben die Art und Weise, wie Informationen verbreitet und bezogen werden, grundlegend verändert. Während früher primär Bücher, Lexika, Zeitschriften und der direkte Austausch mit Peers und Expert*innen genutzt wurden, um gewünschte Auskünfte zu erhalten, und auch genutzt werden mussten, da keine anderen Quellen zur Verfügung standen hat sich der Informationsbezug, in dem im letzten Jahrzehnt vom Analogen hin zum Digitalen bewegt (Bawden & Robinson, 2009). Durch die digitalen Medien steigt die Menge an Inhalten rasant an und sie fungieren als Katalysator für den sogenannten *Information Overload* (Bawden & Robinson, 2020). Dieser beschreibt, dass eine übermäßige Informationsflut nicht mehr hilfreich ist, sondern zur Belastung wird, da sie die Bewältigung von Aufgaben nicht erleichtert, sondern erschwert. Möchte man sich beispielsweise zu einem bestimmten Thema

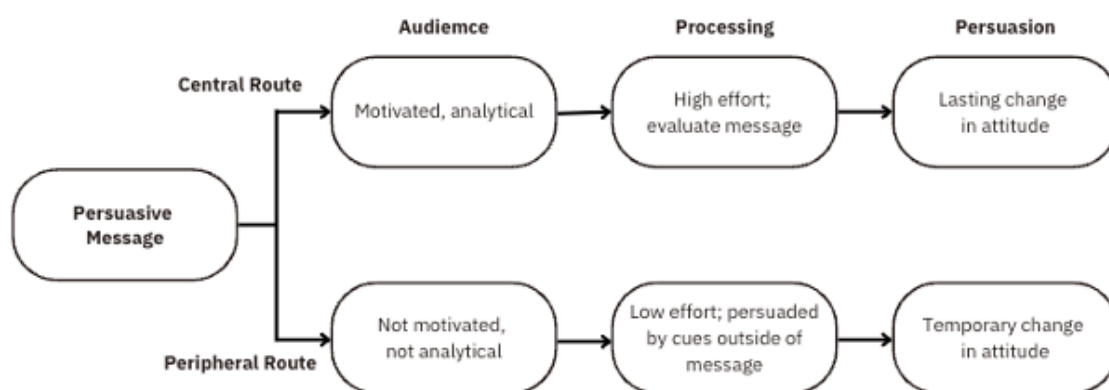
informieren, stößt man über Suchmaschinen oft auf tausende Artikel, Videos und Postings, die sich inhaltlich widersprechen. Dabei ist es nicht immer einfach, die Korrektheit oder Vertrauenswürdigkeit der Quellen zu beurteilen. Das Smartphone spielt in diesem Kontext eine zentrale Rolle: Laut Statista (2024) informieren sich 78,5 % der Studierenden und 53 % der allgemeinen Bevölkerung regelmäßig von unterwegs aus über das Internet mit ihrem Smartphone. Gleichzeitig geben 67,7 % der Studierenden und 50,7 % der Erwachsenen an, dabei stets die gesuchten Informationen zu finden, was das Internet als höchst effektive Datenquelle beschreibt. Allerdings zeigt eine *Critical Online Reasoning Assessment* (CORA) Studie (Depré, 2022), dass Student*innen häufig nicht relevante oder unzulässige Informationen aus dem Internet verwenden. Dabei handle es sich oft um *Fake News* und *Fake Science*.

Der *Fake News* Begriff ist ein noch junger, der als Antwort auf die absichtliche Verbreitung von Falschinformationen im politischen Kontext ab 2016 entstanden ist, da zu dieser Zeit *Fake News* sowohl in der US-Wahl als auch in der Brexit Abstimmung in Großbritannien eine zentrale Rolle spielten (Schwarz & Jalbert, 2020). *Fake News* wird häufig auch als Synonym für *Fake Science* gebraucht. Hierunter fallen erfundene wissenschaftliche Artikel und Ergebnisse, die vor allem online verbreitet werden und teilweise in hochrangigen Fachzeitschriften reproduziert werden (Schwarz & Jalbert, 2020). Gerade in sozialen Netzwerken verbreiten sich *Fake News* und *Fake Science* besonders schnell. Bei einer Analyse von Nachrichten, die auf Twitter geteilt wurden, zeigte sich, dass gerade einmal 41 % der geteilten Inhalte vorher geöffnet wurden, während die übrigen 58 % ausschließlich wegen ihrer Schlagzeile geteilt wurden, ohne dass der Artikel zuvor gelesen wurde (Gabiolkov et al., 2016). Schwarz und Jalbert (2020) erklären dieses Verhalten mit einer intuitiven Reaktion auf das Gesehene: Ist eine Nachricht leicht zu verarbeiten, wird sie eher akzeptiert, während eine tiefere inhaltliche Prüfung ausbleibt. Sie berufen sich dabei

unter anderem auf das *Elaboration Likelihood Model of Persuasion* von Petty und Cacioppo (1986) welches besagt, dass persuasive Informationen je nachdem, wie hoch die *Elaboration Likelihood* ist, entweder über den zentralen Weg oder den peripheren Weg verarbeitet werden (siehe Abbilung 1), Die *Elaboration Likelihood* hängt dabei sowohl von kognitiven Fähigkeiten als auch von der individuellen Motivation ab. Der zentrale Weg ist eine detaillierte Analyse, während der periphere Weg oberflächlich ist und einfache Hinweise wie Attraktivität oder die Quelle bei der Beurteilung berücksichtigt (Petty & Cacioppo, 1984).

Abbildung 1

Elaboration Likelihood Model Petty & Cacioppo (1984)



Ein weiterer theoretischer Ansatz ist das *Reflective-Impulsive Model* von Strack und Deutsch (2004), das zwischen einem langsamen, reflektiven und einem schnellen, impulsiven Verarbeitungssystem unterscheidet. In schnellen, digitalen Kontexten, wie beim Scrollen durch soziale Medien oder beim Lesen von Chatbot-Antworten, könnte das impulsive System aktiv sein, da es auf einfache, leicht verarbeitbare Reize reagiert. Flüssig formulierte KI-Inhalte könnten aus diesem Grund eher akzeptiert werden, ohne dass sie systematisch hinterfragt werden.

Das Zeitalter der Chatbots

Mit der Veröffentlichung des frei zugänglichen Chatbots ChatGPT (OpenAI, 2025) Ende 2022 entstand eine weitere Disruption im Informationssektor. Der Chatbot kann als potenzielle Strategie zur Bewältigung des durch das Internet verstärkten *Information Overload* (Bawden & Robinson, 2020) betrachtet werden, da er Informationen strukturiert, vereinfacht und verdichtet bereitstellt.

Insbesondere im Kontext von Smartphones und der ständigen Verfügbarkeit von Informationen auf mobilen Geräten (Statista, 2024), ermöglicht ChatGPT eine schnelle und einfache Informationsaufnahme, ohne die Nutzer*innen mit übermäßigen Details zu überfordern. Die Antworten beruhen dabei auf Wahrscheinlichkeiten und ersetzen somit ein manchmal als mühsam empfundenenes Screening von Suchmaschinentreffern, da die verfügbaren Quellen in integrierter Form dargeboten werden (Borji, 2023).

Diese schnelle und einfache Wissensvermittlung ist ein klarer Vorteil KI-gestützter Systeme, wodurch sie im alltäglichen Leben einen bedeutenden Mehrwert bieten können, Darstellung, könnte ein besonders hohes Maß an Vertrauen begünstigen. Vertrauen wird auch als die Neigung verstanden, ein Risiko einzugehen – in der Überzeugung, dass das Ergebnis wahrscheinlich positiv sein wird (Glikson & Woolley, 2020; Hoff & Bashir, 2015; Rousseau et al, 1998). Diese Art der Vertrauensbildung kann jedoch dazu führen, dass Ergebnisse unkritisch übernommen werden, selbst wenn sie Fehler enthalten (Borji, 2023). Beim Vertrauen in KI-Systeme spielt insbesondere die wahrgenommene Akkuratheit eine zentrale Rolle. Personen, welche ein hohes Maß an Vertrauen gegenüber KI-gestützten Systemen aufbringen, bewerten diese als zuverlässig und präzise (Molina & Sundar, 2022).

Gleichzeitig sind Chatbots aber wie Blackboxes: Der genaue Vorgang wie eine Antwort generiert wird, bleibt für die Benutzer*innen im Verborgenen und ist somit nicht vollständig nachvollziehbar. Vertrauen in die Richtigkeit der Antworten spielt somit eine übergeordnete Rolle im Umgang mit KI-generierten Inhalten. Wie von Eschenbach (2021)

beschreibt, ist mangelnde Transparenz ein zentrales Problem, welches das Vertrauen in KI-Systeme untergräbt. Nutzer*innen können zwar die Eingaben und Ausgaben eines Systems sehen, beurteilen und gegebenenfalls überprüfen, nicht aber die zugrunde liegenden Entscheidungsprozesse (Eiband et al., 2018). Besonders im Kontext von Chatbots scheint es sinnvoll, dass nicht nur Antworten bereitgestellt, sondern auch deren Entstehung erklärt wird. Hier setzt *Explainable Artificial Intelligence* an. Laut Zednik (2019) zielt dieser Ansatz nicht nur darauf ab, KI-Systeme nachträglich erklärbar zu machen, sondern auch darauf, unterschiedliche Stakeholder*innen und deren spezifische Erklärungsbedarfe zu berücksichtigen. Während Entwickler*innen wissen müssen, wie ein Modell funktioniert, benötigen Nutzer*innen vor allem verständliche Begründungen für Entscheidungen.

Um Nutzer*innen besser vor irreführenden oder fehlerhaften Inhalten zu schützen, könnten *Explainable Artificial Intelligence*-Prinzipien auch auf Chatbots wie ChatGPT angewendet werden und den Arbeitsprozess visualisieren. Es reicht nicht aus, nur Antworten bereitzustellen – ebenso entscheidend ist es, die Entstehung dieser Antworten transparent zu machen, um einen kompetenten Umgang mit Chatbots zu fördern.

Diese Notwendigkeit wird besonders relevant, da moderne KI-Systeme wie ChatGPT zum Halluzinieren neigen, also zur Erzeugung falscher Informationen auf Basis sprachlicher Wahrscheinlichkeiten (Bender et al., 2021; Yoo et al., 2024). Dieser Effekt, auch stochastisches Nachplappern genannt, kann dazu führen, dass falsche Erklärungen als ebenso vertrauenswürdig wahrgenommen werden wie korrekte (Danry et al., 2022). Da es bisher keine klare Lösung für dieses Problem gibt, besteht ein dringender Forschungsbedarf zur menschlichen Wahrnehmung und Reaktion auf KI-generierte Fehlinformationen (Yoo et al., 2024).

Psychologische Einflussfaktoren im Interaktionsdesign von Chatbots

KI-gestützte Chatbots wie ChatGPT werden zunehmend als Informationsquelle genutzt. Aufgrund ihrer klaren, kohärenten Ausdrucksweise vermitteln sie einen hohen Eindruck von Verlässlichkeit, obwohl sie auf Wahrscheinlichkeiten basieren und keine faktenbasierten Aussagen treffen können (Borji, 2023). Daraus ergibt sich die Gefahr, dass fehlerhafte oder erfundene Inhalte ungeprüft übernommen werden.

Unterschiedliche psychologische Mechanismen tragen dazu bei, dass Nutzer*innen solchen Aussagen Vertrauen schenken. So postulierte Spinoza (1677/1982), dass der Akt des Verstehens mit einem automatischen Glaubensprozess einhergeht, während die Ablehnung einer Information zusätzlichen kognitiven Aufwand erfordert. Gilbert (1991) konnte diesen Prozess empirisch nachweisen: Menschen glauben zunächst das, was sie verstehen, und überprüfen Inhalte meist nur bei starken Widersprüchen. Im Kontext überzeugender KI-Antworten ist das besonders kritisch, da formale Stimmigkeit mit inhaltlicher Richtigkeit verwechselt werden kann (Stanley, Whitehead, & Marsh, 2022; Kiziler, 2024).

Auch der *Automation Bias* könnte dabei eine Rolle spielen, Nutzer*innen tendieren laut diesem Bias dazu, Empfehlungen automatisierter Systeme als verlässlicher einzustufen als menschliche Quellen (Goddard, Roudsari, & Wyatt, 2012). Cecil et al. (2023) konnten in einer Studie zeigen, dass auch in einem sensiblen Bereich wie der Personalauswahl fehlerhafte KI-Empfehlungen übernommen wurden, vor allem, wenn sie plausibel präsentiert waren.

Zusätzlich trägt der Fluency-Effekt dazu bei, dass sprachlich flüssige, gut lesbare Aussagen als überzeugender und wahrheitsgetreuer empfunden werden als weniger leserliche (Reber & Schwarz, 1999). Informationen, die leichter verarbeitet werden können, werden positiver bewertet – unabhängig vom Wahrheitsgehalt. Dieser Effekt ist besonders relevant für Chatbots, die genau auf hohe sprachliche Flüssigkeit optimiert sind. Alter und

Oppenheimer (2009) betonen, dass diese Leichtigkeit zu geringerer kognitiver Tiefe und reduzierter kritischer Reflexion führt.

Ein weiterer relevanter Mechanismus, der an der Verarbeitungsflüssigkeit ansetzt, ist der Truth-Effekt. Dieser beschreibt, dass wiederholt gehörte oder vertraut wirkende Aussagen eher als wahr empfunden werden (Renner, 2016). Da Chatbots auf bekannte Sprachmuster zurückgreifen, können sie diesen Effekt verstärken. Nutzer*innen bewerten Aussagen möglicherweise nicht aufgrund ihres Inhalts, sondern aufgrund sprachlicher Vertrautheit als korrekt (Borji, 2023).

Danry et al. (2024) beobachtete, dass das Vertrauen in die Richtigkeit von Chatbots generierte Fehlinformationen erhöht wurde, wenn diese mit logisch wirkenden Erklärungen untermauert waren. Solche trügerischen Erklärungen verstärken den Glauben an falsche Informationen. Ob eine Erklärung hinterfragt wird, hängt stark davon ab, wie logisch sie erscheint – je weniger plausibel, desto eher wird sie als falsch erkannt (Danry et al., 2024).

Ähnliche Effekte sind bereits aus der zwischenmenschlichen Kommunikation bekannt (Rozenblit & Keil 2002) Menschen neigen dazu, oberflächlichen, aber kohärent klingenden Erklärungen Glauben zu schenken, selbst, wenn diese kaum Informationsgehalt besitzen. Rozenblit und Keil (2002) zeigten mit ihrer Forschung zur *Illusion of Explanatory Depth*, dass Menschen oft fälschlicherweise glauben, komplexe Sachverhalte besser zu verstehen, als es tatsächlich der Fall ist. Erst wenn sie gebeten werden, eine detaillierte Erklärung abzugeben, wird ihnen bewusst, wie lückenhaft ihr Wissen ist. Diese Täuschung tritt besonders bei erklärenden Informationen auf, weniger jedoch bei der Erinnerung an Fakten oder Abläufe.

Da Chatbots besonders gut darin sind, flüssige und strukturierte Antworten zu generieren, können sie diese Effekte verstärken. Nutzer*innen nehmen gut formulierte, aber inhaltlich unzureichende Erklärungen eher als wahr an, selbst wenn sie keine fundierten

Informationen liefern. Wie Danry et al. (2024) zeigten, kann dies dazu führen, dass Menschen falschen Informationen aus Chatbot-Antworten mit scheinbar logischen Erklärungen stärker vertrauen als korrekten, aber weniger überzeugend präsentierten Aussagen.

Um die Gefahr von Fehlinformationen zu reduzieren, setzen KI-Anbieter wie OpenAI auf Warnhinweise („ChatGPT kann Fehler machen. Überprüfe wichtige Informationen.“, OpenAI, 2025). Doch Studien legen nahe, dass solche Warnungen oft nicht die beabsichtigte Wirkung haben. Beispielsweise beobachteten Skurnik et al., (2005), dass Warnhinweise nicht nur ineffektiv sein können, sondern sogar ins Gegenteil umschlagen können. Informationen, die mit einer Warnung versehen wurden, bleiben oft besonders gut im Gedächtnis haften und langfristig steigt dadurch die Wahrscheinlichkeit, dass diese Informationen später als glaubwürdig eingestuft werden (Skurnik et al., 2005).

Weitere Forschung bestätigt sich damit, dass nicht alle Warnhinweise gleichermaßen wirksam sind. Kiziler (2024) testete drei verschiedene Bedingungen: keine Warnung, den ChatGPT-Warnhinweis und einen Furchtappell mit. Der Furchtappell führte dabei zu einer signifikant höheren Erkennungs- und Korrekturrate von KI-Halluzinationen im Vergleich zu einer fehlenden Warnung. Zimmermann (2025) beobachtete zudem, dass integrale Fehlerhinweise, also direkt in die Chat-Konversation integrierte Hinweise, die Überprüfung von Antworten fördern können, während explizite Warnhinweise das Nutzer*innenverhalten nicht signifikant zu beeinflussen scheinen. Gleichzeitig beobachtete sie, dass digitale Reife mit dem Vertrauen in den Chatbot in Zusammenhang steht: Personen mit einem höheren Grad an digitaler Reife vertrauten KI-generierten Antworten weniger stark.

Unsicherheitsmarker und ihre Wirkung

Es stellt sich die Frage, wie Nutzer*innen dazu angeregt werden können, generierte Inhalte kritisch zu hinterfragen. Eine Möglichkeit besteht in der Verwendung sprachlicher

Unsicherheitsmarker, die explizit auf Fehler oder eine geringe Wahrscheinlichkeit der Richtigkeit der Antworten hinweisen. Psychologische Forschungen zeigen, dass rhetorische Mittel die Verarbeitung von Informationen beeinflussen können. Klare und verständliche Botschaften werden tendenziell als glaubwürdiger wahrgenommen (Reber & Schwarz, 1999). Im Kontext KI-generierter Inhalte könnte eine weniger technische, weniger prägnante Sprache dazu führen, dass Nutzer*innen die Informationen mit größerer Skepsis betrachten und eher überprüfen. Technische und präzise formulierte Antworten werden hingegen häufiger bevorzugt und als glaubwürdiger eingeschätzt – unabhängig davon, ob sie tatsächlich korrekt sind oder nicht. Dies unterstreicht, dass rhetorische Elemente in der Sprache von KI-Systemen die Wahrnehmung von Wahrheit erheblich verzerren können (Yoo et al., 2024).

Aktuelle Studien zeigen zudem, dass nicht nur das Vorhandensein, sondern auch die Gestaltung von Unsicherheitsmarkern entscheidend ist. Kim, Hwang und Lee (2024) fanden heraus, dass personalisierte Unsicherheitsformulierungen wie „Ich bin mir nicht sicher, aber ...“ die kritische Auseinandersetzung mit der Information fördern und das Vertrauen in die Antwort reduzieren. Allgemeinere Formulierungen wie „Es ist unklar, aber ...“ zeigten hingegen einen geringeren Effekt. Dies legt nahe, dass die Art der Formulierung darüber entscheidet, ob Unsicherheitsmarker tatsächlich zur Reflexion anregen.

Die genannten Befunde deuten darauf hin, dass sprachliche Gestaltung, insbesondere konkrete Unsicherheitsformulierungen wie der Konjunktiv, eine größere Wirkung auf die Erkennung fehlerhafter Inhalte haben könnte als allgemeine Warnhinweise. Hypothese 1: Fehler in KI-generierten Antworten werden eher erkannt, wenn sie im Konjunktiv formuliert sind, als wenn sie nur mit einem allgemeinen Warnhinweis versehen werden.

Die Rolle der Digitalen Reife im Kontext des der Chatbot Nutzung

Ein essenzieller Bestandteil des Informationsbezugs, insbesondere im digitalen Kontext, bleibt auch im Zeitalter von Chatbots der verantwortungsvolle und kritische Umgang mit Inhalten. Es ist zudem erwähnenswert, dass Studien nahelegen, dass Computerspezialist*innen weniger Vertrauen in Künstliche Intelligenz haben als der Durchschnitt, da sie deren Funktionsweise besser verstehen und sich der Unzulänglichkeiten der Technologie besonders bewusst sind (Molina & Sundar, 2022). Diese Erkenntnis bestätigt die Bedeutung der *digitalen Reife* in diesem Kontext, da sie eine fundierte Auseinandersetzung mit den Grenzen und Potenzialen digitaler Technologien ermöglicht (van Swol et al., 2025).

Digitale Reife beschreibt die Fähigkeiten und Einstellungen, die Individuen benötigen, um digitale Technologien auf eine Weise zu nutzen, die sowohl ihr persönliches Wachstum als auch ihre gesellschaftliche Integration unterstützt und ursprünglich speziell für den Kinder- und Jugendkontext entwickelt (Laaber et al., 2023). Die digitale Reife basiert auf zwei wichtigen Aspekten: dem Wachstum und der Anpassung an digitale Kontexte, welche durch Konzepte der psychosozialen Reife und Selbstbestimmung beeinflusst werden (Bleidorn, 2015; Deci & Ryan, 2000; Laaber et al., 2023). Besonders relevant ist dabei die Fähigkeit zur autonomen Nutzung von digitalen Technologien, die es jungen Menschen ermöglicht, diese in Übereinstimmung mit ihren eigenen Zielen und Werten zu verwenden (Livingstone, Mascheroni, G., & Stoilova, 2021; Laaber et al., 2023). Diese Perspektive trägt dazu bei, die Nutzung von digitalen Technologien als einen aktiven und selbstbestimmten Prozess zu verstehen, der über die bloße Vermeidung von problematischem Verhalten hinausgeht und zur Förderung des individuellen Wachstums und einer positiven sozialen Integration beiträgt (Laaber et al., 2023). Laaber und Kolleg*innen (2023) haben zur Messung der *digitalen Reife* das *Digital Maturity Inventory* (DIMI) für Kinder und Jugendliche entwickelt. Das Instrument wurde darüber hinaus in Studien mit erwachsenen

Teilnehmer*innen angewendet, unter anderem bei Zimmermann (2025) im Kontext von KI-Halluzinationen. Dabei konnte gezeigt werden, dass *digitale Reife* negativ mit der Fehlerreplikation in KI-generierten Antworten korrelierte. Personen mit einem hohen digitalen Reifegrad überprüften KI-generierte Lösungen häufiger. Hypothese 2: Der Grad an digitaler Reife beeinflusst, wie oft fehlerhafte KI-Antworten übernommen werden – je höher die digitale Reife, desto seltener werden Fehler weiterverwendet.

Da sprachliche Gestaltung und *digitale Reife* unterschiedliche, aber potenziell kombinierte Einflussfaktoren darstellen, ist davon auszugehen, dass ihre Effekte nicht unabhängig voneinander wirken.

Hypothese 3: Es besteht ein Interaktionseffekt zwischen Formulierungsart (Konjunktiv vs. Warnhinweis) und digitaler Reife auf die Erkennung und Weiterverwendung fehlerhafter Inhalte.

Die Studie

Aufbauend auf diesen Erkenntnissen untersucht die vorliegende Studie, wie sprachliche Unsicherheitsmarker hier in Form des Konjunktivs das Vertrauen in Chatbot-Antworten beeinflussen. Während Yoo et al. (2024) generell die Wirkung sprachlicher Gestaltung auf die Wahrnehmung betrachten, legt diese Arbeit den Fokus darauf, systematisch zu vergleichen, inwiefern Unsicherheitskommunikation die Fehlererkennung bei Nutzer*innen beeinflusst. Dies ist besonders relevant vor dem Hintergrund bestehender Modelle zur Informationsverarbeitung, wie dem *Reflective-Impulsive Model* (Strack & Deutsch, 2004) und dem *Elaboration Likelihood Model* (Petty & Cacioppo, 1986), die erklären, unter welchen Bedingungen Menschen Informationen kritisch reflektieren oder unbewusst übernehmen. Zudem soll untersucht werden, inwiefern dieser Umgang mit der *digitalen Reife* (Laaber et al., 2023) zusammenhängt – also der Fähigkeit, digitale Medien kritisch und kompetent zu nutzen.

Aus den bisherigen theoretischen Überlegungen und den Ergebnissen bestehender Studien ergibt sich folgende Forschungsfrage: *Inwiefern beeinflussen sprachliche Unsicherheitsmarker und die digitale Reife die Erkennung und Weiterverwendung potenziell falscher KI-generierter Inhalte?*

Diese Forschungsfrage soll mithilfe von drei zentralen Hypothesen beantwortet werden: Hypothese 1: Fehler in KI-generierten Antworten werden eher erkannt, wenn sie im Konjunktiv formuliert sind, als wenn sie nur mit einem allgemeinen Warnhinweis versehen werden. Hypothese 2: Der Grad an digitaler Reife beeinflusst, wie oft fehlerhafte KI-Antworten übernommen werden – je höher die digitale Reife, desto seltener werden Fehler weiterverwendet. Hypothese 3: Es besteht ein Interaktionseffekt zwischen Formulierungsart (Konjunktiv vs. Warnhinweis) und digitaler Reife auf die Erkennung und Weiterverwendung fehlerhafter Inhalte.

Diese Untersuchung trägt nicht nur zur theoretischen Debatte über die Gestaltung vertrauenswürdiger KI-Kommunikation bei, sondern soll auch praxisnahe Möglichkeiten aufzeigen, Chatbots so zu gestalten, dass Fehlinformationen reduziert werden und eine kritischere Verarbeitung von Inhalten gefördert wird.

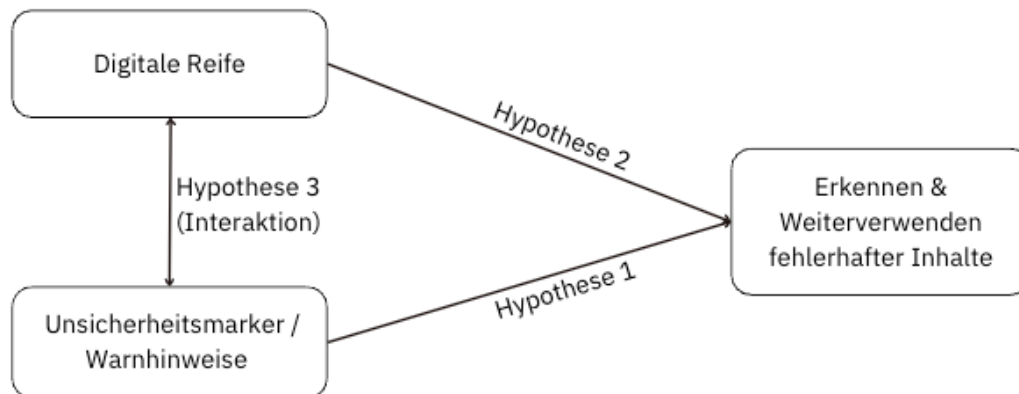
Operationalisierung der Variablen

In dieser Studie gibt es zwei unabhängige Variablen, zum einen die Unsicherheitsmarker (Versuchsbedingung) und zum anderen Standardwarnhinweise (Kontrollbedingung). Die Unsicherheitsmarker in Form des Konjunktivs sind in der Hälfte der Antworten enthalten, während die andere Hälfte den Warnhinweis enthält. Die abhängige Variable ist die Vertrauensbewertung der Chatbot-Antworten nach jedem Stimulus und welches Vorgehen bei der Lösung gewählt wurde. Zusätzlich wurde die *digitale Reife* als moderierende Variable erfasst. Sie beschreibt, wie kompetent jemand im Umgang mit

digitalen Informationen ist und ob diese Ausprägung den Einfluss von Unsicherheitsmarkern (Versuchsbedingung) oder Warnhinweisen (Kontrollbedingung) verändert.

Abbildung 2

Forschungsmodell



Methodik

Stichprobenumfang und Power-Analyse

Zur Bestimmung der erforderlichen Stichprobengröße wurde eine a-priori Power-Analyse durchgeführt. Da für Linear Mixed Models bislang keine etablierten Standardformeln für die Power-Berechnung existieren, orientierte sich die Analyse an einer vergleichbaren vereinfachten Modellstruktur (Varianzanalyse mit Messwiederholung, zwei Messzeitpunkte, ein Between-Subject-Faktor). Unter Annahme eines mittleren Effekts (Cohen's $f = .25$), $\alpha = .05$ und einer Power von .80 ergab die Berechnung eine erforderliche Stichprobengröße von 67 Teilnehmenden ($N = 67$). Diese Planung gewährleistet, dass ein mittlerer Effekt mit einer Wahrscheinlichkeit von 80 % bei einem Signifikanzniveau von .05 entdeckt werden kann.

Teilnehmende

Die Stichprobe bestand aus 69 Proband*innen, davon wurden 94,2 % ($n = 65$) über das LAB-System der Universität Wien rekrutiert. Dieses System ermöglicht es Personen, die an psychologischen Studien teilnehmen möchten, sich zu registrieren, während Forschende der Universität Wien Stichproben für ihre Studien ziehen können. Eingeladen wurden Psychologiestudierende im Bachelorstudium, die für ihre Teilnahme drei LABS-Credits zur Anrechnung für ein Lehrveranstaltung erhielten. Die übrigen 5,8 % ($n = 4$) wurden über eine Studierenden WhatsApp-Gruppe eingeladen und nahmen gegen eine Aufwandsentschädigung von 10 € teil. Die Teilnahme dauerte im Durchschnitt 43 Minuten. Die Teilnehmer*innen waren zwischen 18 und 28 Jahren alt, wobei das Medianalter bei 21 Jahren lag. In der Stichprobe waren 72,5 % weibliche, 24,6 % männliche und 2,9 % diverse Personen. Die Mehrheit der Teilnehmenden stammte aus Deutschland (44,9 %) oder Österreich (43,5 %), während 2,9 % aus der Schweiz, 10,1 % aus anderen EU-Staaten und 1,4 % aus Nicht-EU-Staaten kamen. Der höchste Bildungsabschluss von 82,6 % der Teilnehmenden war die Hochschulreife (Abitur/Matura) und 17,4 % verfügten über einen Studienabschluss.

Design

Es wurde ein Within-Between-Design eingesetzt. Die sprachliche Gestaltung der Chatbot-Antworten (Konjunktiv vs. allgemeiner Warnhinweis) wurde innerhalb der Teilnehmenden variiert (Within-Subject-Faktor), sodass jede Person beide Bedingungen sah. Die Reihenfolge, in der die Antwortarten präsentiert wurden, wurde zwischen den Teilnehmenden variiert (Between-Subject-Faktor), um Reihenfolgeeffekte zu kontrollieren. Die Zuweisung erfolgte randomisiert.

Die Teilnehmenden wurden per Zufallszuteilung mithilfe des Zufallsfilters in zwei Gruppen eingeteilt: Gruppe „Konjunktiv“ sah zunächst zehn Antworten von Chatbot 1 (Versuchsbedingung: Konjunktiv), gefolgt von zehn Antworten von Chatbot 2 (Kontrollbedingung: Warnhinweis), während Gruppe „Kontroll“ die umgekehrte Reihenfolge

durchlief. Auf diese Weise konnte geprüft werden, ob die Reihenfolge der Präsentation einen Einfluss auf das Vertrauen oder das Antwortverhalten hatte.

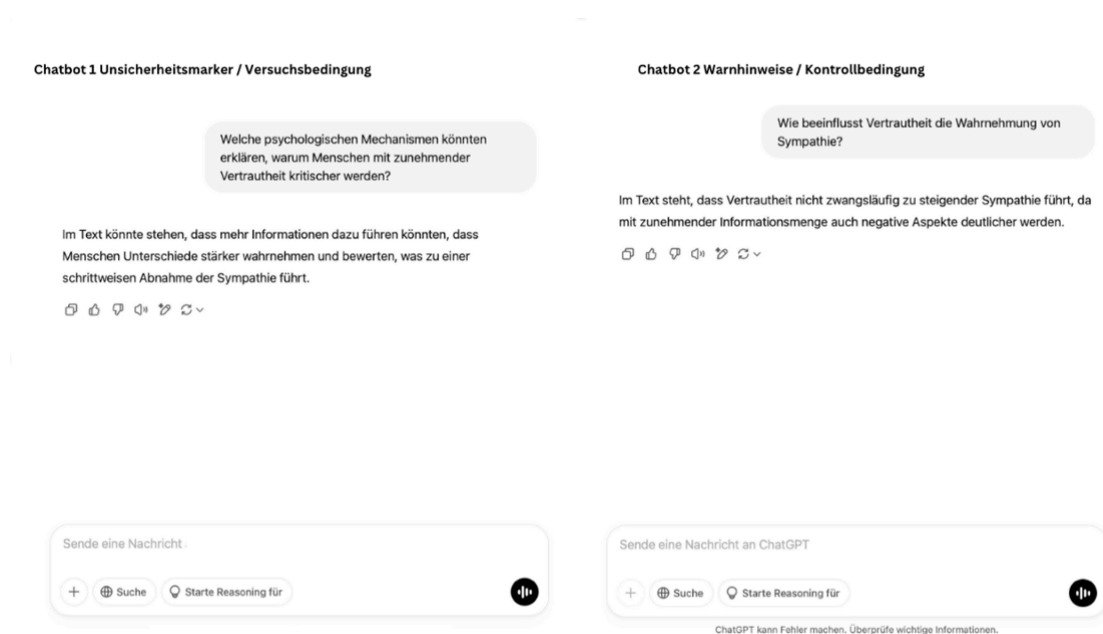
Prozedur

Die Datenerhebung wurde im Frühjahr 2025 über die Plattform Unipark durchgeführt. Nach einem kurzen Briefing durch die Versuchsleitung wurden die Teilnehmenden über die Datenschutzrichtlinien der Universität Wien sowie die Studienbedingungen informiert. Es wurde klargestellt, dass sie die Antworten der Chatbots in dem angegebenen Zeitrahmen beurteilen sollten, ohne jede einzelne Antwort im Originaltext überprüfen zu können. Die Eingaben in das offene Antwortfeld sollten in vereinfachter Form erfolgen.

Im Hauptteil des Fragebogens erhielten alle Teilnehmenden 20 Chatbot-Antworten (siehe Abbildung 3) zu einem einheitlichen Text – dem wissenschaftlichen Artikel *Less is more: The lure of ambiguity, or why familiarity breeds contempt* von Norton, Frost und Ariely (2007). Für jeweils zehn Antworten pro Chatbot galt: Fünf Antworten waren korrekt, fünf enthielten bewusst eingebaute Fehler.

Abbildung 3

Stimuli



Nach jeder Antwort mussten die Teilnehmenden die richtige Antwort selbst eintragen (zur Kontrolle der Fehlererkennung), die Vertrauenswürdigkeit der Antwort auf einer fünfstufigen Skala einschätzen (Vertrauensbewertung) und angeben, ob sie die Antwort übernehmen oder im Paper nachlesen wollten (Vorgehen).

Nach jeweils zehn Antworten wurde das Vertrauen in den jeweiligen Chatbot mit dem *Trust in Automation Questionnaire* (Körber, 2019) erhoben. Dieses Instrument basiert auf einem theoretischen Rahmenmodell, das Vertrauen in automatisierte Systeme als mehrdimensionales Konstrukt versteht. Der Fragebogen erfasst unter anderem Aspekte wie die Leistungsfähigkeit des Systems (z. B. Zuverlässigkeit), die Absicht hinter dem Systemdesign sowie die Nachvollziehbarkeit der Systemprozesse. Insgesamt enthält der Fragebogen 19 Items (siehe Anhang A), die auf einer siebenstufigen Skala bewertet werden, und ermöglicht damit eine differenzierte Einschätzung des Vertrauens in verschiedene Facetten von Automatisierung.

Nachdem beide Versuchsabschnitte abgeschlossen waren, gaben die Teilnehmenden zusätzlich eine Vertrauenspräferenz an, indem sie auf die Frage „Welchem der beiden Chatbots vertrauen Sie mehr?“ antworteten. Diese Einschätzung diente dazu, neben den standardisierten Skalenwerten auch eine bewusste subjektive Entscheidung zwischen den beiden Chatbot-Formaten zu erfassen.

Zum Abschluss des Fragebogens wurde die digitale Reife der Teilnehmenden erfasst. Hierzu wurde das *Digital Maturity Inventory* (Laaber et al., 2023) verwendet. Dieses Instrument misst digitale Reife in den beiden übergeordneten Dimensionen Wachstum durch digitale Technologien und Anpassung an digitale Kontexte. Es besteht aus 32 Items (siehe Anhang B), welche unterschiedliche Aspekte wie digitale Kompetenzen, autonome Nutzung, Selbstregulation in digitalen Räumen sowie soziale Verantwortung online abbilden. Die

Beantwortung erfolgt auf einer fünfstufigen Likert-Skala und ermöglicht so eine umfassende Einschätzung der individuellen digitalen Reife.

Statistische Analysen

Alle statistischen Analysen wurden mit IBM SPSS Statistics (Version 30; IBM Corp., 2023) für macOS durchgeführt. Zunächst wurden die *Trust in Automation Questionnaire* (Körber, 2019) -Werte zu einem Score für *Vertrauen in Kompetenz und Reliabilität* sowie *Vertrauen in Automatisierung* aggregiert und zu einer Gesamtbewertung des Vertrauens in einen Chatbot zusammengefasst. Darüber hinaus wurden die zehn Dimensionen der *digitalen Reife* zusammengefasst und anschließend ein gewichteter Gesamtscore gebildet. Zur Vorbereitung der Daten wurden die Variablen für die Vertrauensbewertung der einzelnen Antworten, Gesamtbewertung des Vertrauens in einen Chatbot und das Vorgehen für jede Versuchsperson in das Long-Format überführt, sodass jeweils separate Werte für Chatbot 1 (Versuchsbedingung) und Chatbot 2 (Kontrollbedingung) vorlagen.

Die abhängigen Variablen umfassten (1) die Vertrauensbewertung, also die Einschätzung der Vertrauenswürdigkeit jeder einzelnen Chatbot-Antwort, (2) das Vorgehen, also die Entscheidung, ob die Chatbot-Antwort akzeptiert oder im Originalpaper nachgelesen wurde, sowie (3) die Gesamtbewertung des Vertrauens in einen Chatbot, bestehend aus den Subskalen *Vertrauen in Kompetenz und Reliabilität* und *Vertrauen in Automatisierung*.

Für die Hypothesentests wurden Pearson-Korrelationen und ein Linear Mixed Model (LMM) mit Bedingung (within-subject), Reihenfolge (between-subject) und digitaler Reife (Kovariat) durchgeführt.

Ergebnisse

Korrelationen

Die digitale Reife zeigte weder in der Konjunktiv-Bedingung ($r = -.06$, $p = .483$) noch in der Warnhinweis-Bedingung ($r = -.07$, $p = .416$) einen signifikanten Zusammenhang mit der Bewertung des Chatbots. Hingegen bestand in beiden Bedingungen ein signifikanter negativer Zusammenhang zwischen der Bewertung und dem Vorgehen (Konjunktiv-Bedingung: $r = -.46$, $p < .001$; Warnhinweis-Bedingung: $r = -.50$, $p < .001$), sodass höhere Bewertungen mit seltenerem Nachschlagen in der Originalquelle einhergingen. Hypothese 2, die einen negativen Zusammenhang zwischen digitaler Reife und Weiterverwendung fehlerhafter Inhalte annahm, wurde nicht bestätigt.

Tabelle 1

Korrelationen

		Konjunktiv Bewertung	Kontroll Bewertung	Digitale Reife	Vorgehen
Konjunktiv Bewertung	Pearson_Korrelation	1	.547**	-.060	-.462**
	Sig. (2-seitig)		<.001	.483	<.001
	N	138	138	138	138
Kontroll Bewertung	Pearson_Korrelation	.547**	1	-.070	-.501**
	Sig. (2-seitig)	<.011		.416	<.001
	N	138	138	138	138
Digitale Reife	Pearson_Korrelation	-.060	-.070	1	-.005
	Sig. (2-seitig)	.483	.416		.957
	N	138	138	138	138
Vorgehen	Pearson_Korrelation	-.462**	-.501**	-.005	1
	Sig. (2-seitig)	<.001			
	N	138	138	138	138

Anmerkung: Die Korrelationen basieren auf 138 Beobachtungen (69 Teilnehmende mit je zwei Messungen)

Linear Mixed Model

Zur Überprüfung von Hypothese 1, wonach Fehler in KI-generierten Antworten eher erkannt werden, wenn diese im Konjunktiv formuliert sind, wurde ein lineares gemischtes Modell (LMM) berechnet. Die Bedingung (Konjunktiv vs. Warnhinweis) wurde als Within-Subject-Faktor, die Darbietungsreihenfolge als Between-Subject-Faktor und digitale Reife als Kovariate in das Modell aufgenommen.

Die Analyse ergab einen signifikanten Haupteffekt der Bedingung, $F(1, 33) = 7.21, p = .011, d = 0.46$. Chatbot-Antworten mit Konjunktivformulierungen ($M = 2.61, SE = 0.10$) wurden signifikant kritischer bewertet als Antworten mit allgemeinem Warnhinweis ($M = 2.88, SE = 0.10$). Weder die Reihenfolge ($F(1, 33) = 0.29, p = .594$) noch die *digitale Reife* ($F(1, 33) = 0.29, p = .600$) zeigten signifikante Haupteffekte. Auch die Interaktionen zwischen Bedingung und Reihenfolge ($F(1, 33) = 0.16, p = .690$) sowie zwischen Bedingung und *digitaler Reife* ($F(1, 33) = 0.10, p = .753$) waren nicht signifikant, womit Hypothese 3 (Interaktionseffekt) nicht bestätigt wurde.

Tabelle 2

Ergebnisse des linearen gemischten Modells zur Vorhersage der Bewertung

Prädiktor	F	df1	df2	p	Effektgröße (d)
Bedingung	7.21	1	33	.011	0.46
Reihenfolge_between	0.29	1	33	.594	—
Digitale Reife	0.29	1	33	.600	—
Bedingung × Reihenfolge_between	0.16	1	33	.690	—
Bedingung × Digitale Reife	0.10	1	33	.753	—

Anmerkung: Das LMM basierte auf 138 Beobachtungen (69 Teilnehmende mit je zwei Messungen)

Damit bestätigt sich der Effekt der sprachlichen Gestaltung der Chatbot-Antworten auch nach Kontrolle für die Reihenfolge und die *digitale Reife* der Teilnehmenden. Entgegen der Erwartung zeigte sich jedoch kein Effekt der Bedingung auf die Wahl des Vorgehens (Übernahme oder Nachlesen), was bedeutet, dass die kritischere Bewertung nicht unmittelbar mit einer erhöhten Bereitschaft zur Überprüfung oder Ablehnung der Chatbot-Antworten einherging.

Tabelle 3

Geschätzte Marginalmittel (SE) der Bewertung nach Bedingung

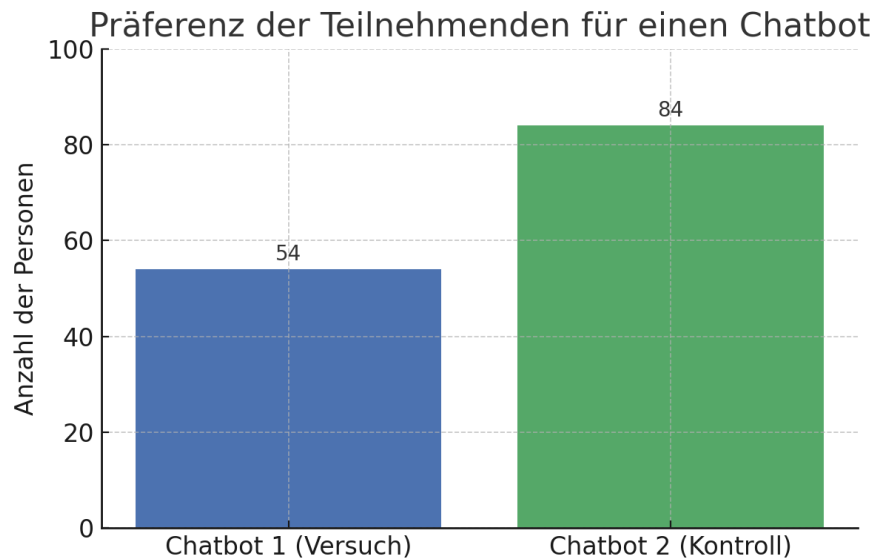
Bedingung	Mittelwert	Std. Fehler	df	95% Konfidenzintervall
Konjunktiv	2.61	.11	49.85	[2.39, 2.83]
Warnhinweis	2.88	.11	49.85	[2.66, 3.01]

Anmerkung: Das LMM basierte auf 138 Beobachtungen (69 Teilnehmende mit je zwei Messungen)

Zusätzlich wurden die Teilnehmenden gebeten, eine Präferenz zwischen den beiden Chatbots anzugeben. 60,9 % ($n = 84$) entschieden sich für Chatbot 2 (Kontrollbedingung), 39,1 % ($n = 54$) für Chatbot 1 (Versuchsbedingung; siehe Abbildung 4). Diese Präferenz zugunsten des Kontroll-Chatbots mit explizitem Warnhinweis steht im Einklang mit den vorherigen Befunden zu einer höheren wahrgenommenen Vertrauenswürdigkeit dieser Form der Fehlerkennzeichnung.

Abbildung 4

Präferenzverteilung der Teilnehmenden hinsichtlich der beiden Chatbots.



Diskussion

Ziel dieser Studie war es, zu untersuchen, wie sprachliche Gestaltung in Form von sprachlichen Unsicherheitsmarkern sowie die *digitale Reife* von Nutzer*innen das Vertrauen in Chatbot-Antworten und den Umgang mit fehlerhaften Inhalten beeinflussen kann. Die Ergebnisse stützen Hypothese 1: Konjunktivformulierungen führten zu einer signifikant kritischeren Bewertung der Antworten als allgemeine Warnhinweise. Diese sprachliche Markierung scheint Nutzer*innen stärker für potenzielle Fehler zu sensibilisieren, führte jedoch nicht zu einer Veränderung des tatsächlichen Verhaltens.

Hypothese 2 konnte nicht bestätigt werden: Es zeigte sich kein signifikanter Zusammenhang zwischen digitaler Reife und Vertrauen in die Chatbots. Die Korrelationen wiesen zwar in beiden Bedingungen in eine negative Richtung, waren jedoch nicht signifikant. Auch Hypothese 3 konnte nicht bestätigt werden: Es ergab sich kein signifikanter Interaktionseffekt zwischen Formulierungsart und digitaler Reife auf die Erkennung oder Weiterverwendung fehlerhafter Inhalte.

Die Befunde stehen im Einklang mit bisherigen Studien zur Vertrauensbildung in KI-Systeme, wonach Transparenz und verständliche Kommunikation eine zentrale Rolle spielen

(von Eschenbach, 2021). Die stärkere Wirkung des Konjunktivs gegenüber eines expliziten Warnhinweis könnte mit dem *Fluency-Effekt* erklärt werden: Komplexere sprachliche Formulierungen (wie der Konjunktiv) führen zu einer umständlicheren Verarbeitung, wodurch Inhalte kritischer hinterfragt werden (Reber, Schwarz & Winkielman, 2004). Bezüglich der Warnhinweise sind die Ergebnisse kohärent mit anderen Studien (Jalbert & Schwarz, 2021), die zeigen konnten, dass Warnhinweise nicht zwangsläufig zu einem kritischeren Umgang mit Informationen führen. Ebenso ist das höhere Vertrauen in den Chatbot, der einen Warnhinweis verwendet, im Einklang mit den Ergebnissen von Zimmermann (2025) und Kiziler (2024), die beide herausfanden, dass der standardisierte Warnhinweis von ChatGPT – der auch in der vorliegenden Studie verwendet wurde – keinen signifikanten negativen Effekt auf das Vertrauen oder das Verhalten der Proband*innen zeigt.

Interessant ist zudem, dass das tatsächliche Verhalten – also die Entscheidung, ob eine Antwort überprüft wurde – durch die sprachliche Gestaltung nicht signifikant beeinflusst wurde. Dies könnte daran liegen, dass Nutzer*innen unter Zeitdruck oder bei Informationsüberflutung auf heuristische Strategien zurückgreifen (Bawden & Robinson, 2009), etwa durch oberflächliche Plausibilitätsurteile. Auch aus Sicht des *Elaboration Likelihood Model* (Petty & Cacioppo, 1986) lässt sich dies interpretieren: Wenn sich Teilnehmende in einem Modus peripherer Verarbeitung befinden – etwa aufgrund begrenzter Motivation oder kognitiver Ressourcen –, dann reichen sprachliche Marker wie der Konjunktiv möglicherweise nicht aus, um eine tiefere inhaltliche Auseinandersetzung anzustoßen. In solchen Fällen wirken Hinweise zwar auf der Bewertungsebene, führen jedoch nicht zwangsläufig zu einer bewussten Entscheidung bezüglich einer Verhaltensänderung – wie etwa dem Nachlesen im Originaltext.

Praktische Implikationen

Für die Gestaltung vertrauenswürdiger KI-Systeme ergeben sich mehrere konkrete Empfehlungen: Konjunktivformulierungen scheinen ein effektives Mittel zu sein, um die kritische Auseinandersetzung mit KI-generierten Inhalten zu fördern. Insbesondere, da sie subtil wirken und weniger aufdringlich sind als plakative Warnhinweise. Gleichzeitig zeigen die Ergebnisse, dass digitale Reife kein einheitlicher Schutzfaktor ist, sondern nur in bestimmten Bereichen greift. Daraus lassen sich Ansätze für individualisierte Unterstützungsangebote oder erklärende Interfaces ableiten, die sich an verschiedenen Nutzer*innentypen orientieren.

Limitationen

Die Studie weist mehrere Limitationen auf: Die Stichprobe war nicht repräsentativ, was die Generalisierbarkeit einschränkt. Auffällig ist, dass es sich, der Art der Rekrutierung geschuldet, um eine junge, akademische und weiblich dominierte Stichprobe handelt. Somit sind die Ergebnisse nicht repräsentativ für die allgemeine österreichische Bevölkerung und nicht ohne weiteres auf andere Kulturen übertragbar.

Zudem könnte ein Lerneffekt aufgetreten sein, da die Teilnehmenden jeweils zehn Antworten pro Chatbot bewerten mussten. Die Konzentration auf ein einziges Paper (Norton et al., 2007) begrenzt außerdem die thematische Vielfalt. Auch der Wert des Vorgehens blieb unbeeinflusst – möglicherweise, weil die Studie auf einem einzigen Paper basierte, was die Notwendigkeit zur Überprüfung reduzierte.

Ein methodischer Aspekt betrifft die Versuchsbedingungen: Da in beiden Bedingungen ein Hinweis auf potenzielle Fehler vorhanden war (entweder in Form eines Konjunktivs oder eines Warnhinweises), gab es keine klassische Kontrollgruppe ohne jeglichen Hinweis. Der Warnhinweis könnte daher ebenfalls eine Auswirkung gehabt haben, was die Interpretation der Effekte erschwert.

Dass die *digitale Reife* nur teilweise Effekte zeigte, lässt sich möglicherweise auf die homogene Stichprobe zurückführen. Die Teilnehmenden waren primär Psychologiestudierende mit hoher digitaler Vorbildung, was die Varianz in den DIMI-Scores einschränken könnte. Schließlich erfolgte die Bewertung der Chatbot-Antworten isoliert und nicht im realen Anwendungskontext, was die externe Validität reduziert. Ein weiterer kritischer Aspekt betrifft die kulturelle Perspektive: Die Studie wurde ausschließlich mit deutschsprachigen Studierenden durchgeführt, was mögliche kulturelle Einflüsse auf das Vertrauensverhalten unberücksichtigt lässt. Frühere Forschung zeigt, dass kultureller Hintergrund eine bedeutende Rolle bei der Einstellung gegenüber KI-Systemen oder Robotern spielt. So fanden Bartneck et al. (2006), dass insbesondere US-amerikanische Teilnehmende deutlich positivere Einstellungen gegenüber Robotern zeigten als japanische oder europäische. Es ist daher denkbar, dass auch das Vertrauen in Chatbots und die Wirksamkeit von Hinweisen je nach kulturellem Kontext unterschiedlich ausfallen.

Ausblick

Zukünftige Forschung sollte verschiedene Antwortformate und Themenbereiche einbeziehen, um die Generalisierbarkeit der Befunde zu erhöhen. Durch den Einsatz einer Kombination sprachlicher und visueller Unsicherheitsmarker – etwa durch Icons, Farbgebung oder Layoutänderungen – könnten differenzierte Wirkungen auf Wahrnehmung, Vertrauen und Verhalten analysiert werden. Denkbar wäre, da die Antworten selbst auf Wahrscheinlichkeitsberechnungen beruhen auch eine Wahrscheinlichkeitsangabe.

Darüber hinaus gäben vertiefende qualitative Verfahren wie Eye-Tracking und Interviews spannende Hinweise, um zu verstehen wie Nutzer*innen KI-Antworten wahrnehmen, bewerten und darauf reagieren. Auf diese Weise könnten die kognitiven Prozesse besser verstanden werden.

Ein weiterer relevanter Aspekt zukünftiger Studien besteht in der Variation der Anwendungskontexte – etwa durch realistischere Settings (z. B. medizinische, rechtliche oder bildungsbezogene Chatbots), um die externe Validität zu erhöhen. Dabei wäre es auch interessant, die alltägliche Nutzung von Chatbots zu untersuchen, anstatt sich nur auf den Umgang mit vorgegebenen Themen zu konzentrieren. Ebenso könnten Langzeitstudien durchgeführt werden, um mögliche Lerneffekte im Umgang mit Unsicherheitsmarkern oder Veränderungen im Vertrauen über die Zeit hinweg zu erfassen.

Schließlich wäre es vielversprechend, den Einfluss kultureller Unterschiede systematisch zu untersuchen, beispielsweise durch internationale Vergleichsstudien. Dabei könnten auch Persönlichkeitsmerkmale, technologische Vorerfahrungen oder kognitive Stile als moderierende Variablen einbezogen werden.

Literatur

- Alter, A. L., & Oppenheimer, D. M. (2009). Uniting the tribes of fluency to form a metacognitive nation. *Personality and Social Psychology Review*, 13(3), 219–235.
<https://doi.org/10.1177/1088868309341564>
- Bawden, D., & Robinson, L. (2009). The dark side of information: Overload, anxiety and other paradoxes and pathologies. *Journal of Information Science*, 35(2), 180–191.
<https://doi.org/10.1177/0165551508095781>
- Bawden, D., & Robinson, L. (2020, 30. Juni). Information overload: An introduction. In *Oxford Research Encyclopedia Politics*.
<https://oxfordre.com/politics/view/10.1093/acrefore/9780190228637.001.0001/acrefore-9780190228637-e-1360>
- Bartneck, C., Suzuki, T., Kanda, T., & Nomura, T. (2006). The influence of people's culture and prior experiences with Aibo on their attitude towards robots. *AI & SOCIETY*, 21(1–2), 217–230. <https://doi.org/10.1007/s00146-006-0052-7>
- 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*, 610–623. ACM.
<https://doi.org/10.1145/3442188.3445922>
- Bhattacharyya, M., Miller, V. M., Bhattacharyya, D., & Miller, L. E. (2023). High Rates of Fabricated and Inaccurate References in ChatGPT-Generated Medical Content. *Cureus*, 15(5), e39238. <https://doi.org/10.7759/cureus.39238>
- Bleidorn, W. (2015). What accounts for personality maturation in early adulthood? *Current Directions in Psychological Science*, 24(3), 245–252. <https://doi.org/10.1177/0963721414568662>
- Borji, Ali. (2023). A Categorical Archive of ChatGPT Failures.

<https://doi.org/10.21203/rs.3.rs-2895792/v1>

Cecil, J., Lerner, E., Hudecek, M., Sauer, J., & Gaube, S. (2023). *The effect of AI-generated advice on decision-making in personnel selection.*

<https://doi.org/10.31219/osf.io/349xe>

Danry, V., Pataranutaporn, P., Epstein, Z., Groh, M., & Maes, P. (2022). Deceptive AI systems that give explanations are just as convincing as honest AI systems in human-machine decision making. *arXiv preprint*, arXiv:2210.08960.

<https://doi.org/10.48550/arXiv.2210.0896>

Danry, V., Chan, J., Stuhlmüller, A., Shah, R., Evans, J., & Steinhardt, J. (2024). Deceptive AI systems that give explanations are more convincing than honest AI systems and can amplify belief in misinformation. *arXiv*.

<https://doi.org/10.48550/arXiv.2408.00024>

Deci, E. L., & Ryan, R. M. (2000). The “what” and “why” of goal pursuits: Human needs and the self-determination of behavior. *Psychological Inquiry*, 11(4), 227–268.

https://doi.org/10.1207/S15327965PLI1104_01

Depré, K. (2022). Critical online reasoning – Validierung neu entwickelter Aufgaben zur Erfassung und Förderung des kritischen Umgangs mit Online-Medien. Springer VS.

<https://doi.org/10.1007/978-3-658-39698-5>

Eiband, M., Schneider, H., Bilandzic, M., Fazekas-Con, C., & Hussmann, H. (2018).

Bringing transparency design into practice. Proceedings of the 23rd International Conference on Intelligent User Interfaces, 211–223.

<https://doi.org/10.1145/3172944.3172961>

von Eschenbach, W. J. (2021). Transparency and the Black Box Problem: Why We Do Not Trust AI. *Philosophy & Technology*, 34(4), 1607–1622.

<https://doi.org/10.1007/s13347-021-00477-0>

- Gabielkov, M., Ramachandran, A., Chaintreau, A., & Legout, A. (2016). Social clicks: What and who gets read on Twitter? *ACM SIGMETRICS Performance Evaluation Review*, 44(1), 179–192. <https://doi.org/10.1145/2964791.2901462>
- Gilbert, D. T. (1991). How mental systems believe. *American Psychologist*, 46(2), 107–119. <https://doi.org/10.1037/0003-066X.46.2.107>
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660. <https://doi.org/10.5465/annals.2018.0057>
- Goddard, K., Roudsari, A., & Wyatt, J. C. (2012). Automation bias: A systematic review of frequency, effect mediators, and mitigators. *Journal of the American Medical Informatics Association*, 19(1), 121–127. <https://doi.org/10.1136/amiajnl-2011-000089>
- Haan, K. (2023, Juli 20). Artificial intelligence and consumer sentiment: Over 75% of consumers are concerned about misinformation. *Forbes Advisor*. Abgerufen am 2. Dezember 2024, von <https://www.forbes.com/advisor/business/artificial-intelligence-consumer-sentiment/>
- Hoff, K. A., & Bashir, M. (2015). *Trust in automation: Integrating empirical evidence on factors that influence trust*. *Human Factors*, 57(3), 407–434. <https://doi.org/10.1177/0018720814547570>
- IBM Corp. (2023). IBM SPSS Statistics for Mac (Version 30) [Computer software]. <https://www.ibm.com/products/spss-statistics>
- Kim, M., Hwang, K., & Lee, J. (2024). *Uncertainty Communication in AI Systems: Effects of First-Person and General Markers on User Trust and Verification*. arXiv preprint arXiv:2405.00623. <https://doi.org/10.48550/arXiv.2405.00623>

- Kiziler, Y. (2024). *KI-Halluzinationen: Wie helfen Warnsignale?* [Master's thesis, Universität Wien].
- Körber, M. (2019). Theoretical considerations and development of a questionnaire to measure trust in automation. In S. Bagnara, R. Tartaglia, S. Albolino, T. Alexander, & Y. Fujita (Eds.), *Proceedings of the 20th Congress of the International Ergonomics Association (IEA 2018)* (Advances in Intelligent Systems and Computing, Vol. 823). Springer. https://doi.org/10.1007/978-3-319-96074-6_2
- Laaber, F., Florack, A., Koch, T., & Hubert, M. (2023). Digital maturity: Development and validation of the Digital Maturity Inventory (DIMI). *Computers in Human Behavior*, 143, 107709. <https://doi.org/10.1016/j.chb.2023.107709>
- Livingstone, S., Mascheroni, G., & Stoilova, M. (2021). The outcomes of gaining digital Skills for young people's lives and well-being: A systematic evidence review (pp. 1–27). *New Media & Society*. <https://doi.org/10.1177/14614448211043189>
- Molina, M. D., & Sundar, S. S. (2022). Does distrust in humans predict greater trust in AI? Role of individual differences in user responses to content moderation. *New Media & Society*. <https://doi.org/10.1177/14614448221103534>
- Norton, M. I., Frost, J. H., & Ariely, D. (2007). Less is more: the lure of ambiguity, or why familiarity breeds contempt. *Journal of personality and social psychology*, 92(1), 97–105. <https://doi.org/10.1037/0022-3514.92.1.97>
- Petty, R. E., & Cacioppo, J. T. (1984). Source factors and the elaboration likelihood model of persuasion. *Advances in Consumer Research*, 11, 668-672. In T. C. Kinnear (Ed.), *Advances in Consumer Research* (Vol. 11, pp. 668-672). Association for Consumer Research.
- Petty, R. E., & Cacioppo, J. T. (1986). The elaboration likelihood model of persuasion.

- In *Communication and persuasion: Central and peripheral routes to attitude change* (pp. 1–24). Springer. https://doi.org/10.1007/978-1-4612-4964-1_1
- Reber, R., & Schwarz, N. (1999). Effects of perceptual fluency on judgments of truth. *Consciousness and Cognition*, 8(3), 338–342. <https://doi.org/10.1006/ccog.1999.0386>
- Reber, R., Schwarz, N., & Winkielman, P. (2004). Processing fluency and aesthetic pleasure: Is beauty in the perceiver's processing experience? *Personality and Social Psychology Review*, 8(4), 364–382. https://doi.org/10.1207/s15327957pspr0804_3
- Renner, C. H. (2016). Validity effect. In R. F. Pohl (Ed.), *Cognitive illusions: Intriguing phenomena in judgement, thinking and memory* (pp. 242–255). Psychology Press. <https://doi.org/10.4324/9781315696935>
- Rousseau, D. M., Sitkin, S. B., Burt, R. S., & Camerer, C. (1998). Not so different after all: A cross-discipline view of trust. *Academy of Management Review*, 23(3), 393–404. <https://doi.org/10.5465/AMR.1998.926617>
- Rozenblit, L.R. and Keil, F.C. (2002) The misunderstood limits of folk science: an illusion of explanatory depth. *Cogn. Sci.* 26, 521–562. https://doi.org/10.1207/s15516709cog2605_1
- Schwarz, N., & Jalbert, M. (2020). WHEN (FAKE) NEWS FEELS TRUE: Intuitions of truth and the acceptance and correction of misinformation. In R. Greifeneder, M. Jaffe, E. Newman, & N. Schwarz (Eds.), *The psychology of fake news* (pp. 73–89). Routledge. <https://doi.org/10.4324/9780429295379>
- Skurnik, I., Yoon, C., Park, D. C., & Schwarz, N. (2005). How warnings about false claims become recommendations. *Journal of Consumer Research*, 31(4), 713–724. <https://doi.org/10.1086/426605>
- Spinoza, B. (1982). *The Ethics and selected letters* (S. Feldman, Ed. and S. Shirley, Trans.). Indianapolis, IN: Hackett. (Original work published 1677)

- Stanley, M. L., Whitehead, P. S., & Marsh, E. J. (2022). The cognitive processes underlying false beliefs. *Journal of Consumer Psychology*, 32(2), 359-369.
<https://doi.org/10.1002/jcpy.1289>
- Strack, F., & Deutsch, R. (2004). Reflective and Impulsive Determinants of Social Behavior. *Personality and Social Psychology Review*, 8(3), 220-247.
https://doi.org/10.1207/s15327957pspr0803_1
- Statista. (2024). *Umfrage unter Studenten in Deutschland zum Informationsbedürfnis*. Statista. Abgerufen am 12. Februar 2025, von
<https://de.statista.com/statistik/daten/studie/865188/umfrage/umfrage-unter-studenten-in-deutschland-zum-informationsbeduerfnis>
- Teo, Y. S. (2024, March 1). What is Interaction Design?. Interaction Design Foundation – IxDF. <https://www.interaction-design.org/literature/article/what-is-interaction-design>
- Unipark. (2025). Umgang mit Chatbots [Umfrage]. <https://www.unipark.com/>
- Van Swol, L., Florack, A., & Reimer, T. (2025). AI as the next frontier in communication and social cognition: artificial communication. In T. Reimer, L. van Swol., & A. Florack (Eds.), *The Routledge handbook of communication and social cognition*. New York, NY: Routledge/Taylor and Francis.
- Yoo, D., Kang, H., & Oh, C. (2024). Deciphering Deception: How Different Rhetoric of AI Language Impacts Users' Sense of Truth in LLMs. *International Journal of Human–Computer Interaction*, 1–21. <https://doi.org/10.1080/10447318.2024.2316370>
- Zednik, C. (2021). Solving the Black Box Problem: A Normative Framework for Explainable Artificial Intelligence. *Philosophy & Technology*, 34(2), 265–288.
<https://doi.org/10.1007/s13347-019-00382-7>
- Zimmermann, I. (2024), *Vertrauen in Künstliche Intelligenz: Die Rolle von digitaler Reife, Unsicherheiten und Warnsignalen* [Master's thesis, Universität Wien].

Anhang

Anhang A

Skalen und Items des Trust in Automation Questionnaire (TiA; Körber, 2019)

#	Item	Subskala
<i>Reliabilität/Kompetenz</i>		
1	R/K1	Das System ist imstande, Situationen richtig einzuschätzen
6	R/K2	Das System arbeitet zuverlässig
10	R/K3*	Ein Ausfall des Systems ist wahrscheinlich
13	R/K4	Das System kann wirklich komplizierte Aufgaben übernehmen
15	R/K5*	Das System könnte stellenweise einen Fehler machen
19	R/K6	Ich bin überzeugt von den Fähigkeiten des Systems
<i>Verständlichkeit/Vorhersagbarkeit</i>		
2	V/V1	Mir war durchgehend klar, in welchem Zustand sich das System befindet
7	V/V2*	Das System reagiert unvorhersehbar
11	V/V3	Ich konnte nachvollziehen, warum etwas passiert ist
16	V/V4*	Zu erkennen, was das System als nächstes macht, ist schwer
<i>Vertrautheit</i>		
3	Ve1	Ich kenne bereits ähnliche Systeme
17	Ve2	Ich habe ähnliche Systeme bereits genutzt
<i>Intention der Entwickler</i>		
4	I1	Die Entwickler sind vertrauenswürdig
8	I2	Die Entwickler nehmen mein Wohlergehen ernst
<i>Neigung zu vertrauen</i>		
5	N1*	Bei unbekannten automatisierten Systemen sollte man eher vorsichtig sein

12	N2	Ich vertraue einem System eher als dass ich ihm misstraue
18	N3	Automatisierte Systeme funktionieren generell gut
		<i>Vertrauen in Automation</i>
9	ViA1	Ich vertraue dem System
14	ViA2	Ich kann mich auf das System verlassen

Anmerkung. * invertiert formuliert. Adaptiert übernommen von Körber (2019).

Anhang B

Digital Maturity Inventory (DIMI; Laaber, Florack, Koch & Hubert, 2023)

Variablenname	Item
<i>Autonomy Within</i>	
<i>Digital Contexts</i>	<i>Wenn ich ein Mobilgerät benutze...</i>
autonomy_c1	... bin ich online, weil ich denke, immer online sein zu müssen
autonomy_c2	... bin ich online, weil ich sonst das Gefühl habe, irgendetwas zu verpassen
autonomy_c3	... habe ich das Gefühl, dass es mein Leben bestimmt
autonomy_w1	... suche ich mir die Inhalte aus, die ich sehen möchte
autonomy_w2	... tue ich das, was mir selbst Spaß macht
autonomy_w3	... entscheide ich selbst, was ich tue
$\alpha = .70$, CR = .70, AVE = .44 (N=558)	
<i>Digital Literacy</i>	
literacy_1	Ich weiß, wie ich die Privatsphäre-Einstellungen anpassen kann (z.B. Cookies deaktivieren)
literacy_2	Ich weiß, wie man die Privatsphäre-Einstellungen in sozialen Medien anpassen kann (z.B. Instagram, Snapchat, TikTok)
literacy_3	Ich weiß, wie ich Fotos, Dokumente, oder andere Dateien in einer Cloud speichern kann (z.B. Google Drive, iCloud)
$\alpha = .82$, CR = .82, AVE = .61 (N=558)	

Individual Growth

in Digital Contexts Wenn ich ein Mobilgerät benutze...

- | | |
|----------|--------------------------------|
| growth_1 | ... lerne ich neue Dinge |
| growth_2 | ... lerne ich nützliche Dinge |
| growth_3 | ... lerne ich neue Fähigkeiten |

$\alpha = .81$, CR = .81, AVE = .59 (N=558)

Digital Risk

Awareness Wenn ich ein Mobilgerät benutze..."

- | | |
|--------|--|
| risk_1 | ... bin ich sehr vorsichtig |
| risk_2 | ... ist mir meine eigene Sicherheit sehr wichtig |
| risk_3 | ... achte ich drauf, vorsichtig zu sein |
-

Support-Seeking Regarding Digital Problems

- | | |
|-----------|--|
| support_1 | ... frage ich andere um Hilfe, wenn ich ein Problem habe |
| support_2 | ... und ich nicht weiterweiß, dann hole ich mir Hilfe von anderen |
| support_3 | ... und ich ein technisches Problem habe, frage ich andere um Hilfe (z.B. einen Freund/eine Freundin, Elternteil, Geschwister) |
| support_4 | ... und Probleme mit anderen im Internet habe, hole ich mir Hilfe |

$\alpha = .80$, CR = .81, AVE = .52 (N=558)

Anmerkung. * invertiert formuliert. Adaptiert übernommen von Laaber et al., 2023)

Anhang C

Druckversion

21.04.25, 14:51

Fragebogen**1 Startseite****VPN****2 Standardseite****Vielen Dank für Ihr Interesse an meiner Studie!**

Im Rahmen dieser Studie der Universität Wien interessiere ich mich für den Umgang mit KI generierten Inhalten. Ich wäre Ihnen sehr dankbar, wenn Sie sich bereit erklären, die vorliegende wissenschaftliche Untersuchung mit Ihrer wertvollen Studienteilnahme zu unterstützen!

Die Studie dauert etwa 45 Minuten und es ist wichtig, dass Sie die Fragen aufmerksam durchlesen und wahrheitsgemäß beantworten.

Ihre Teilnahme an dieser Studie erfolgt freiwillig. Sie können jederzeit, ohne Angabe von Gründen, Ihre Bereitschaft zur Teilnahme ablehnen oder auch im Verlauf der Studie zurückziehen. Die Ablehnung der Teilnahme oder ein vorzeitiges Ausscheiden aus dieser Studie hat keine nachteiligen Folgen für Sie. Ihre Angaben werden völlig anonym erhoben und können nicht mit Ihnen in Verbindung gebracht werden. Zusätzlich werden Ihre Daten streng vertraulich behandelt. Ihre Daten werden nicht an Dritte weitergegeben oder kommerziell genutzt.

Wenn Sie damit einverstanden sind und mit der Studie beginnen möchten, dann klicken Sie bitte auf "JA".

Pauline Schmidt

☐ JA**3 Demographie**

Druckversion

21.04.25, 14:51

Welches Geschlecht haben Sie?

- ☐ weiblich
- ☐ männlich
- ☐ divers
- ☐ keine Angabe

Wie alt sind Sie?

Bitte geben Sie Ihr Alter in Ziffern an (z.B. 32).

Welche Staatsangehörigkeit haben Sie?

- ☐ Deutschland
- ☐ Österreich
- ☐ Schweiz
- ☐ anderer EU-Staat
- ☐ anderer Nicht-EU-Staat

Welchen höchsten allgemeinbildenden Schulabschluss haben Sie?

- ☐ Ohne Abschluss
- ☐ Pflichtschulabschluss
- ☐ Realschulabschluss oder einen vergleichbaren Abschluss
- ☐ Fachschule oder berufsbildenden Schule
- ☐ Hochschulreife / Abitur / Matura
- ☐ Studienabschluss
- ☐ Berufsausbildung
- ☐ Ich habe einen anderen Schulabschluss und zwar

4 Instruktion

Bitte stellen Sie sich vor, Sie möchten Fragen zu einem wissenschaftlichen Paper beantworten.

Insgesamt haben Sie maximal 30 Minuten Zeit, um 20 Fragen zu beantworten, die sich ausschließlich auf den Inhalt des Papers beziehen. Aus ökonomischen Gründen wurde die Studie und die einzelnen Fragen in zwei verschiedene Chatbots hochgeladen. Sie werden insgesamt 20 Fragen beantworten: 10 Fragen werden Ihnen im ersten Chatbot angezeigt und 10 im zweiten Chatbot. Beide Chatbots basieren auf demselben Large Language Model, unterscheiden sich jedoch in der Darstellung.

Nach jeder Frage sehen Sie einen Screenshot des jeweiligen Chatbots. Sie haben die Möglichkeit, entweder die Antwort des Chatbots zu überprüfen oder ihr direkt zu vertrauen und sie zu übernehmen. Wenn Sie sie überprüfen möchten, dann kontrollieren Sie die Antwort indem Sie in dem geöffneten Paper von Norton und Forst (2007) nachlesen.

Nach jeder Frage werden Sie aufgefordert, die richtige Antwort in ein Textfeld einzugeben und die Vertrauenswürdigkeit der Antwort zu bewerten.

Anschließend werden noch einige Fragen zu Ihrem Umgang mit digitalen Technologien gestellt.

Sie können sich Ihre Lösungen notieren, um zu prüfen wie oft Sie richtig lagen und erhalten nach der Studie die Auflösung.

Wenn Sie bereit sind, drücken Sie auf „Weiter“, um die erste Frage zu erhalten.

Viel Erfolg!

5.1 Konjunktiv

Druckversion

21.04.25, 14:51

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1

2

3

4

5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?☐ Ich habe die Antwort des Chatbots übernommen☐ Ich habe im Paper nachgelesen**5.1.3 Standardseite**

Welche psychologischen Mechanismen könnten erklären, warum Menschen mit zunehmender Vertrautheit kritischer werden?

Im Text könnte stehen, dass mehr Informationen dazu führen könnten, dass Menschen Unterschiede stärker wahrnehmen und bewerten, was zu einer schrittweisen Abnahme der Sympathie führt.

📄 👍 🗨️ 🔊 🗑️ ↺

Sende eine Nachricht



Suche

Starte Reasoning für

**Geben Sie die korrekte Antwort hier ein:**

Dieser Block ist von Chatbot 1

In diesem Abschnitt werden Sie Screenshots von Chatbot 1 sehen. Bitte lesen Sie die Antworten des Chatbots sorgfältig und entscheiden Sie, ob Sie ihnen vertrauen oder sie weiter überprüfen möchten, indem Sie im Paper nachlesen. Nach der Beantwortung der Fragen geben Sie die korrekte Antwort in das Textfeld ein und bewerten, wie vertrauenswürdig Sie die Antwort finden.

5.1.1 Standardseite

Wie könnte Ungewissheit genutzt werden, um soziale Bindungen zu fördern?

Im Text könnte stehen, dass gezielt aufrechterhaltene Ungewissheit das Interesse an einer Person steigern kann, weil das Gehirn positive Merkmale eher ergänzt als negative.

< 1/2 >      

Sende eine Nachricht



Suche



Starte Reasoning für



Geben Sie die korrekte Antwort hier ein:

5.1.2 Standardseite

Geben Sie die korrekte Antwort hier ein:

5.1.6 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

5.1.7 Standardseite

5.1.4 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

5.1.5 Standardseite

Was wurde in der Studie über die Entwicklung von Sympathie in langfristigen Beziehungen festgestellt?

Im Text könnte stehen, dass sich Menschen mit der Zeit immer sympathischer werden, da Vertrautheit automatisch zu einer stärkeren Bindung führt.

📄 📌 🔍 🗣️ 🔄 📄

Sende eine Nachricht



Suche



Starte Reasoning für



Druckversion

21.04.25, 14:51

Welche langfristigen Auswirkungen könnte die Wahrnehmung von Ähnlichkeit haben?

Im Text könnte stehen, dass anfangs wahrgenommene Ähnlichkeiten mit der Zeit an Bedeutung verlieren, während Dissimilaritäten stärker ins Bewusstsein rücken.

Sende eine Nachricht

+

Suche

Starte Reasoning für

Geben Sie die korrekte Antwort hier ein:

5.1.8 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1

2

3

4

5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

https://ww3.unipark.de/www/print_survey.php?syid=1059068&__menu_node=print

Seite 8 von 36

5.1.9 Standardseite

Wie beeinflusst Unsicherheit die Wahrnehmung einer Person in sozialen Interaktionen?

Im Text könnte stehen, dass Unsicherheit in sozialen Interaktionen überwiegend negativ bewertet wird, weil Menschen sich instinktiv nach klaren und eindeutigen Informationen sehnen.



Sende eine Nachricht



Suche



Starte Reasoning für



Geben Sie die korrekte Antwort hier ein:

5.1.10 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1

2

3

4

5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

- ☐ Ich habe die Antwort des Chatbots übernommen
- ☐ Ich habe im Paper nachgelesen

5.1.11 Standardseite

Wie beeinflusst das Ausmaß der geteilten
Informationen die zwischenmenschliche Nähe?

Im Text könnte stehen, dass Menschen sich umso enger miteinander verbunden fühlen, je mehr persönliche Informationen sie voneinander wissen, da Transparenz das Vertrauen stärkt.

📄 👍 🗣️ 📄 🔄

Sende eine Nachricht



🌐 Suche

💡 Starte Reasoning für



Geben Sie die korrekte Antwort hier ein:

5.1.12 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?☐ Ich habe die Antwort des Chatbots übernommen☐ Ich habe im Paper nachgelesen**5.1.13 Standardseite**

Wie könnte kultureller Hintergrund die Wahrnehmung von Ungewissheit verändern?

Im Text könnte stehen, dass Menschen aus kollektivistischen Kulturen Ungewissheit positiver bewerten als Menschen aus individualistischen Kulturen, da sie stärker an sozialen Kontexten orientiert sind.

📄 👍 🗣️ 🔍 ↺ ↻

Sende eine Nachricht



Suche



Starte Reasoning für



Geben Sie die korrekte Antwort hier ein:

5.1.14 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

5.1.15 Standardseite

Welche Rolle könnte die Selbstwahrnehmung in der
Einschätzung von anderen spielen?

Im Text könnte stehen, dass Menschen andere positiver bewerten, wenn sie sich
selbst gerade in einem positiven Licht sehen.

📄 👍 👎 🔊 🔍 ↺ ↻

Sende eine Nachricht



🌐 Suche

💡 Starte Reasoning für



Geben Sie die korrekte Antwort hier ein:

5.1.16 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

5.1.17 Standardseite

Druckversion

21.04.25, 14:51

Was besagt die Studie über die langfristigen Auswirkungen von Unsicherheit in einer Freundschaft?

Im Text könnte stehen, dass Unsicherheit in Freundschaften vermieden werden sollte, da sie das Vertrauen schwächt und dazu führt, dass Menschen sich emotional voneinander distanzieren.

📄 📌 🗨️ 🗑️ 🔄 📄

Sende eine Nachricht



🌐 Suche

💡 Starte Reasoning für



Geben Sie die korrekte Antwort hier ein:

5.1.18 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

Druckversion

21.04.25, 14:51

5.1.19 Standardseite

Welche Methode wurde in der dritten Untersuchung zur Messung von Sympathie verwendet?

Im Text könnte stehen, dass die dritte Studie mithilfe der „Emotional Response Scale (ERS)“ Sympathiebewertungen erfasste.

📄 👍 🗨️ 🗑️ ✎️ ↺

Sende eine Nachricht



🌐 Suche

🗨️ Starte Reasoning für



Geben Sie die korrekte Antwort hier ein:

5.1.20 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Das System kann wirklich komplizierte Aufgaben übernehmen

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Das System könnte stellenweise einen Fehler machen

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Ich bin überzeugt von den Fähigkeiten des Systems

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Ich vertraue dem System

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

- ☐ Ich habe die Antwort des Chatbots übernommen
- ☐ Ich habe im Paper nachgelesen

5.1.21 Einschätzung des Chatbots

Bitte geben Sie nun an wie sehr Sie den Aussagen über Chatbot 1 zustimmen.

Das System ist imstande, Situationen richtig einzuschätzen

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Das System arbeitet zuverlässig

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Ein Ausfall des Systems ist wahrscheinlich

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Ich kann mich auf das System verlassen

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

5.2 Standardseite**Dieser Block ist von Chatbot 2**

In diesem Abschnitt werden Sie Screenshots von Chatbot 2 sehen. Bitte lesen Sie die Antworten des Chatbots sorgfältig und entscheiden Sie, ob Sie ihnen vertrauen oder sie weiter überprüfen möchten, indem Sie im Paper nachlesen. Nach der Beantwortung der Fragen geben Sie die korrekte Antwort in das Textfeld ein und bewerten, wie vertrauenswürdig Sie die Antwort finden.

5.2.1 Standardseite

Wie beeinflusst Vertrautheit die Wahrnehmung von Sympathie?

Im Text steht, dass Vertrautheit nicht zwangsläufig zu steigender Sympathie führt, da mit zunehmender Informationsmenge auch negative Aspekte deutlicher werden.

📄 📌 🗑️ 🗑️ 🗑️ 🗑️

Sende eine Nachricht an ChatGPT



Suche

Starte Reasoning für



ChatGPT kann Fehler machen. Überprüfe wichtige Informationen.

Geben Sie die korrekte Antwort hier ein:

5.2.2 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

5.2.3 Standardseite

Welche Rolle spielt der sogenannte „Selbstwahrnehmungs-Bias“ in der Studie?

Im Text steht, dass Menschen sich selbst attraktiver einschätzen, wenn sie weniger Informationen über ihren Datingpartner vor dem Treffen haben.

📄 👍 🗣️ 🔗 🔄

Sende eine Nachricht an ChatGPT

+ 🌐 Suche 💡 Starte Reasoning für



ChatGPT kann Fehler machen. Überprüfe wichtige Informationen.

Geben Sie die korrekte Antwort hier ein:

5.2.4 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Druckversion

21.04.25, 14:51

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

- ☐ Ich habe die Antwort des Chatbots übernommen
- ☐ Ich habe im Paper nachgelesen

5.2.5 Standardseite

Was besagt die Dissimilaritäts-Kaskade über
Langzeitinteraktionen?

Im Text steht, dass wahrgenommene Unterschiede mit der Zeit an Bedeutung gewinnen und die Sympathie verringern können.

📄 👍 🗨️ 🔊 🔄 ⌵

Sende eine Nachricht an ChatGPT



Suche



Starte Reasoning für



ChatGPT kann Fehler machen. Überprüfe wichtige Informationen.

Geben Sie die korrekte Antwort hier ein:

5.2.6 Standardseite

5.2.8 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

5.2.9 Standardseite

Wie wird Unsicherheit im zwischenmenschlichen Kontext beschrieben?

Im Text steht, dass Unsicherheit in sozialen Beziehungen sowohl positive als auch negative Auswirkungen haben kann, abhängig vom Kontext und der individuellen Interpretation.

📄 📌 🗣️ 🎧 🔄 📄

Sende eine Nachricht an ChatGPT



🌐 Suche

💡 Starte Reasoning für



ChatGPT kann Fehler machen. Überprüfe wichtige Informationen.

Druckversion

21.04.25, 14:51

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

5.2.7 Standardseite

Was wurde in der „Studie 2“ experimentell untersucht?

Im Text steht, dass Studienteilnehmer durch das Trinken von Kaffee ihre Sympathie gegenüber unbekannten Personen signifikant steigerten.

📄 📌 🗑️ 🗑️ 🗑️ 🗑️

Sende eine Nachricht an ChatGPT



🌐 Suche

💡 Starte Reasoning für



ChatGPT kann Fehler machen. Überprüfe wichtige Informationen.

Geben Sie die korrekte Antwort hier ein:

Geben Sie die korrekte Antwort hier ein:

5.2.10 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

5.2.11 Standardseite

5.2.13 Standardseite

Welcher Persönlichkeitsfaktor war laut Studie der stärkste Prädiktor für langfristige Sympathie?

Im Text steht, dass emotionale Stabilität als stärkster Prädiktor für eine stabile und positive Langzeitbeziehung identifiziert wurde.

📄 👍 🗨️ 🔊 🔄 ⌵

Sende eine Nachricht an ChatGPT



🌐 Suche

💡 Starte Reasoning für



ChatGPT kann Fehler machen. Überprüfe wichtige Informationen.

Geben Sie die korrekte Antwort hier ein:

5.2.14 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1

2

3

4

5

Druckversion

21.04.25, 14:51

Welche Rolle spielt die Informationsmenge bei der Einschätzung von Personen?

Im Text steht, dass mehr Informationen über eine Person nicht zwangsläufig zu einer positiven Einschätzung führen, da negative Details stärker ins Gewicht fallen können.

📄 🗑️ 🔍 🗨️ 🔄 📌

Sende eine Nachricht an ChatGPT



🌐 Suche

💡 Starte Reasoning für



ChatGPT kann Fehler machen. Überprüfe wichtige Informationen.

Geben Sie die korrekte Antwort hier ein:

5.2.12 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

Druckversion

21.04.25, 14:51

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?☐ Ich habe die Antwort des Chatbots übernommen☐ Ich habe im Paper nachgelesen**5.2.17 Standardseite**

Wie unterscheidet sich die Wahrnehmung von Unsicherheit in neuen und bestehenden Beziehungen?

Im Text steht, dass Unsicherheit in neuen Beziehungen oft als aufregend wahrgenommen wird, während sie in bestehenden Beziehungen als störend empfunden werden kann.

📄 👍 🗣️ 🔍 🔄

Sende eine Nachricht an ChatGPT



Suche



Starte Reasoning für



ChatGPT kann Fehler machen. Überprüfe wichtige Informationen.

Geben Sie die korrekte Antwort hier ein:

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

- ☐ Ich habe die Antwort des Chatbots übernommen
- ☐ Ich habe im Paper nachgelesen

5.2.15 Standardseite

Welche Strichprobe gab es bei Studie 5?

Im Text steht, dass es in Studie 5 210 Teilnehmer*innen gab (N = 210; 102 Männer, 108 Frauen, Alter M = 37,8 Jahre, SD = 11,8).

📄 📌 🔍 🗨️ 🔄 📄

Sende eine Nachricht an ChatGPT

+ 🔍 Suche 🗨️ Starte Reasoning für



ChatGPT kann Fehler machen. Überprüfe wichtige Informationen.

Geben Sie die korrekte Antwort hier ein:

5.2.16 Standardseite

5.2.18 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

5.2.19 Standardseite

Welcher Aspekt wurde in der abschließenden Analyse besonders hervorgehoben?

Im Text steht, dass das Vertrauen zwischen den Probanden durch eine abschließende Reflexionsrunde signifikant gestärkt wurde.

📄 👍 🗣️ 🔊 🖋️ ↺

Sende eine Nachricht an ChatGPT



🌐 Suche

💡 Starte Reasoning für



ChatGPT kann Fehler machen. Überprüfe wichtige Informationen.

Geben Sie die korrekte Antwort hier ein:

5.2.20 Standardseite

Wie vertrauenswürdig war die Chatbot Antwort für Sie?

1 = gar nicht vertrauenswürdig

5 = sehr vertrauenswürdig

1 2 3 4 5

Wie sind Sie bei der Beantwortung der Frage vorgegangen?

☐ Ich habe die Antwort des Chatbots übernommen

☐ Ich habe im Paper nachgelesen

5.2.21 Kopie von Einschätzung des Chatbots

Bitte geben Sie nun an wie sehr Sie den Aussagen über Chatbot 2 zustimmen.

Das System ist imstande, Situationen richtig einzuschätzen

☐ Stimme gar nicht zu

☐ Stimme eher nicht zu

☐ Stimme weder zu noch nicht zu

☐ Stimme eher zu

☐ Stimme voll zu

Das System arbeitet zuverlässig

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Ein Ausfall des Systems ist wahrscheinlich

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Das System kann wirklich komplizierte Aufgaben übernehmen

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Das System könnte stellenweise einen Fehler machen

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Ich bin überzeugt von den Fähigkeiten des Systems

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Ich vertraue dem System

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

Ich kann mich auf das System verlassen

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Stimme weder zu noch nicht zu
- ☐ Stimme eher zu
- ☐ Stimme voll zu

6 Vergleich der Chatbots**Welchem Chatbot vertrauen Sie mehr?**

- ☐ Chatbot 1
- ☐ Chatbot 2

7 Digitale Reife

Vielen Dank, dass Sie den ersten Teil des Fragebogens ausgefüllt haben!

Im nächsten Abschnitt geht es um Ihren Umgang mit digitalen Technologien. Bitte lesen Sie die Fragen aufmerksam durch und beantworten Sie sie ehrlich und gewissenhaft.

Für diesen Teil gibt es keine zeitliche Begrenzung – nehmen Sie sich also gerne die Zeit, die Sie benötigen.

Bitte denken Sie daran, wie oft die folgenden Dinge normal passieren, wenn Sie ein Mobilgerät (z.B. Handy, Tablet, Laptop) benutzen.

Wenn ich ein Mobilgerät benutze...

	Niemals	Selten	Manchmal	Oft	Immer
bin ich online, weil ich denke, immer online sein zu müssen	☐	☐	☐	☐	☐
bin ich online, weil ich sonst das Gefühl habe, irgendetwas zu verpassen	☐	☐	☐	☐	☐
habe ich das Gefühl, dass es mein Leben bestimmt	☐	☐	☐	☐	☐

Wenn ich ein Mobilgerät benutze...

	Niemals	Selten	Manchmal	Oft	Immer
suche ich mir die Inhalte aus, die ich sehen möchte	☐	☐	☐	☐	☐
tue ich das, was mir selbst Spaß macht	☐	☐	☐	☐	☐
entscheide ich selbst, was ich tue	☐	☐	☐	☐	☐

Wenn ich ein Mobilgerät benutze...

	Niemals	Selten	Manchmal	Oft	Immer
lerne ich neue Dinge	☐	☐	☐	☐	☐
lerne ich nützliche Dinge	☐	☐	☐	☐	☐
lerne ich neue Fähigkeiten	☐	☐	☐	☐	☐
kann ich Probleme alleine lösen	☐	☐	☐	☐	☐

Wenn ich ein Mobilgerät benutze...

	Niemals	Selten	Manchmal	Oft	Immer
und mich jemand online oder in einer Nachricht kritisiert, reagiere ich sofort ohne vorher über die Konsequenzen nachzudenken	☐	☐	☐	☐	☐
und mich eine Nachricht wütend macht, reagiere ich zu schnell und bereue es danach	☐	☐	☐	☐	☐
und mich jemand beleidigt, versuche ich es ihm/ihr heimzuzahlen	☐	☐	☐	☐	☐

Wenn ich ein Mobilgerät benutze...

	Niemals	Selten	Manchmal	Oft	Immer
achte ich auf meine Wortwahl, wenn ich mit jemandem nicht einer Meinung bin, um nicht gemein zu wirken	☐	☐	☐	☐	☐
achte ich auf die Gefühle von anderen Personen	☐	☐	☐	☐	☐
und Bilder von anderen Personen poste oder verschicke, achte ich darauf, dass diese niemanden beleidigen oder in Schwierigkeiten bringen	☐	☐	☐	☐	☐
respektiere ich die Meinung und das Wissen von anderen Personen	☐	☐	☐	☐	☐

Wenn ich ein Mobilgerät benutze...

	Niemals	Selten	Manchmal	Oft	Immer
verwende ich es, um das Leben in meiner Umgebung, meiner Stadt oder in der Welt insgesamt zu verbessern	☐	☐	☐	☐	☐
verwende ich das Internet, um Kampagnen, zum Beispiel für Umweltschutz, zu unterstützen oder um auf den Klimawandel mehr aufmerksam zu machen	☐	☐	☐	☐	☐
verwende ich es, um mich für Dinge einzusetzen, die wirklich wichtig in der Welt sind	☐	☐	☐	☐	☐

Wenn ich ein Mobilgerät benutze...

	Niemals	Selten	Manchmal	Oft	Immer
bin ich sehr vorsichtig	☐	☐	☐	☐	☐
ist mir meine eigene Sicherheit sehr wichtig	☐	☐	☐	☐	☐
achte ich darauf, vorsichtig zu sein	☐	☐	☐	☐	☐

Wenn ich ein Mobilgerät benutze...

	Niemals	Selten	Manchmal	Oft	Immer
frage ich andere um Hilfe, wenn ich ein Problem habe	☐	☐	☐	☐	☐
und ich nicht weiterweiß, dann hole ich mir Hilfe von anderen	☐	☐	☐	☐	☐
und ich ein technisches Problem habe, frage ich andere um Hilfe (z.B. einen Freund/eine Freundin, Elternteil, Geschwister)	☐	☐	☐	☐	☐
und Probleme mit anderen im Internet habe, hole ich mir Hilfe	☐	☐	☐	☐	☐

Wenn ich ein Mobilgerät benutze...

	Niemals	Selten	Manchmal	Oft	Immer
und ich online verärgert oder verletzt werde, dauert es lange, bis ich mich wieder besser fühle	☐	☐	☐	☐	☐
und ich online verärgert oder verletzt werde, bleibe ich lange in einer schlechten Stimmung	☐	☐	☐	☐	☐
und ich online verärgert oder verletzt werde, gibt es nur wenige Dinge, die mich wieder aufheitern	☐	☐	☐	☐	☐

Bitte wählen Sie "Niemals" aus...

	Niemals	Selten	Manchmal	Oft	Immer
Bitte wählen Sie "Niemals" aus um zu zeigen, dass Sie den Fragebogen gewissenhaft ausfüllen.	☐	☐	☐	☐	☐

Bitte geben Sie an, wie sehr die folgenden Aussagen auf Sie zutreffen, wenn Sie daran denken, wie Sie das Internet oder Mobilgeräte wie Handys, Tablets oder Laptops nutzen. Stellen Sie sich vor, wie sehr die Aussagen auf Sie zutreffen, wenn Sie die Dinge genau jetzt, alleine machen müssten.

	Gar nicht	Eher nicht	Teils-teils	Ziemlich	Völlig
Ich weiß, wie ich die Privatsphäre-Einstellungen anpassen kann (z.B. Cookies deaktivieren)	☐	☐	☐	☐	☐
Ich weiß, wie man die Privatsphäre-Einstellungen in sozialen Medien anpassen kann (z.B. Instagram, Snapchat, TikTok)	☐	☐	☐	☐	☐
Ich weiß, wie ich Fotos, Dokumente, oder andere Dateien in einer Cloud speichern kann (z.B. Google Drive, iCloud)	☐	☐	☐	☐	☐
Ich weiß, wie man kontrollieren kann, ob die Informationen wahr sind, die man online findet	☐	☐	☐	☐	☐
Ich weiß, wie man herausfinden kann, ob man einer Website vertrauen kann	☐	☐	☐	☐	☐

8 Endseite

Der Fragebogen ist nun abgeschlossen, wenn Sie Interesse an den Ergebnissen haben melden Sie sich gerne per Email a01641505@unet.univie.ac.at

Vielen Dank für Ihre Teilnahme!