



MASTERARBEIT | MASTER'S THESIS

Titel | Title

KI-Influencer vs. menschliche Influencer: Eine vergleichende
Analyse der Werbewirksamkeit auf Instagram

verfasst von | submitted by

Nenad Radakovic BSc (WU)

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 066 915

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

Masterstudium Betriebswirtschaft

Betreut von | Supervisor:

ao. Univ.-Prof. Mag. Dr. Heribert Reisinger

Dankesagung

Ein außerordentlicher Dank gebührt meinem Betreuer ao. Univ.-Prof. Dr. Heribert Reisinger, der mir wertvolle Impulse für diese Arbeit gegeben hat. Insbesondere verdanke ich ihm die Idee, parasoziale Beziehungen und Response-Variablen in meine Untersuchung aufzunehmen, sowie entscheidende Hinweise zur Experimentgestaltung. Seine Unterstützung hat mich stets in meinem Vorgehen bestärkt und motiviert.

Mein besonderer Dank gilt meiner Freundin Ronja, die mich stets motiviert und unterstützt hat – und mit viel Geduld ertragen hat, dass ich oft bis in die frühen Morgenstunden geschrieben habe. Ebenso danke ich von Herzen meiner Schwester Jelena, meiner Mutter Zora und meinem Vater Dragomir für ihre liebevolle und verständnisvolle Unterstützung und ihren Zuspruch, die mir viel Kraft gegeben haben.

Nicht zuletzt danke ich meinem Arbeitgeber, sowie meinen Kunden, die mir während meines Masterstudiums äußerst entgegengekommen sind und Verständnis dafür hatten, dass mein Studium für mich oberste Priorität hatte.

Inhaltsverzeichnis

1. Einleitung	1
2. Theoretischer Hintergrund und Hypothesenbildung	2
2.1 Definition.....	3
2.2 Studienüberblick.....	4
2.3 Wahrnehmung.....	9
2.3.1 Erkennung von KI-Influencern	9
2.3.2 Unheimlichkeit	10
2.3.3 Coolness, Neuheit und Innovation	11
2.3.4 Personalisierung und Kommerzialisierung	11
2.3.5 Nützlichkeit.....	12
2.4 Psychologische Mechanismen	12
2.4.1 Sozialpsychologische Distanz.....	12
2.4.2 Authentizität.....	13
2.4.3 Vertrauen	13
2.4.4 Identifikation	14
2.4.5 Parasoziale Beziehungen	14
2.5 Einstellungen	15
2.5.1 Marken-Einstellung	15
2.5.2 Anzeigen-Einstellung	15
2.6 Kauf- und Verhaltensabsichten	16
2.6.1 Kaufabsicht	16
2.6.2 Influencer-Folgeabsicht	16
2.7 Hypothesen	17
2.8 Forschungsmodell	20
3. Methodik	21
3.1 Stichprobenbeschreibung	21
3.2 Operationalisierung	22
3.3 Experimentdesign	24
3.3.1 Forschungsdesign	24
3.3.2 Stimulusmaterial	24
3.3.3 Prozedur	26
4. Ergebnisse	29
4.1 Variablenkonstruktion und Testvoraussetzungen.....	29
4.2 Deskriptive Statistik	31
4.3 Schließende Statistik.....	37
5. Diskussion	44
Literaturverzeichnis	50

Abstract (Deutsch)

Diese Studie untersucht die Werbewirksamkeit von KI- und menschlichen Influencern auf Instagram in einem 2×2-Between-Subjects-Design mit den Faktoren Influencer-Typ (KI vs. Mensch) und Produktkategorie (Technik vs. Kosmetik). Menschliche Influencer erzielen höhere Anzeigen-Einstellungen, stärkere parasoziale Beziehungen und mehr Kongruenz, während sich die Marken-Folgeabsicht, Marken-Einstellung und Engagement-Rate nicht statistisch signifikant unterscheiden. Parasoziale Beziehungen und Kongruenz vermitteln den Effekt des Influencer-Typs vollständig auf die Anzeigen-Einstellung, die wiederum auf die Marken-Einstellung, Engagement-Rate und Marken-Folgeabsicht wirkt. Eine Moderation durch die Produktkategorie wurde nicht bestätigt. Erstmals wurden parasoziale Beziehungen, Marken-Folgeabsicht und Engagement-Rate bei KI-Influencern in einer quantitativen Studie untersucht.

Abstract (English)

This study investigates the advertising effectiveness of AI and human influencers on Instagram through a 2×2 between-subjects design, focusing on influencer type (AI vs. human) and product category (technology vs. cosmetics). The findings reveal that human influencers outperform AI influencers in terms of ad attitude, parasocial relationships, and congruence. However, no statistically significant differences were observed for brand-following intention, brand attitude, and engagement rate. Parasocial relationships and congruence fully mediated the effect of influencer type on ad attitude, which subsequently influenced brand attitude, engagement rate, and brand-following intention. The study found no moderation effect by product category. Notably, this is the first quantitative study to examine parasocial relationships, brand-following intention, and engagement rate specifically for AI influencers, offering new insights into their role in advertising effectiveness.

Tabellenverzeichnis

Tabelle 1: Überblick über den aktuellen Forschungsstand zu KI-Influencern	5
Tabelle 2: Operationalisierung der abhängigen Variablen und Mediatoren	23
Tabelle 3: Zusammenfassung der deskriptiven Werte	32
Tabelle 4: Korrelationstabelle	37

Abbildungsverzeichnis

Abbildung 1: Instagram-Post: Miquela.....	10
Abbildung 2: Instagram-Post: Lopez	10
Abbildung 3: Forschungsmodell	21
Abbildung 4: Influencer-Post	27
Abbildung 5: Spotify-Post	27
Abbildung 6: Labello-Post	27
Abbildung 7: Ablaufplan des Experiments	29
Abbildung 8: Boxplot Marken-Einstellung	33
Abbildung 9: Boxplot-Anzeigen-Einstellung	34
Abbildung 10: Boxplot Marken-Folgeabsicht	35
Abbildung 11: Boxplot Kongruenz	36
Abbildung 12: Boxplot Parasoziale Beziehung	37

1. Einleitung

Social-Media-Influencer spielen eine wichtige Rolle in Werbekampagnen. Sie können mit ihrer Originalität und Reichweite die Werbeeffektivität steigern und neue Zielgruppen erreichen (Leung et al., 2022). Für 2024 werden weltweit Investitionen von 32 Milliarden Dollar in Influencer-Kooperationen prognostiziert (Statista, 2024). Allerdings beschränkt sich der Einsatz von Influencern längst nicht mehr auf menschliche Akteure. Ein Teil dieser Investitionen fließt in das mit 39 % jährlich wachsende Segment der computergenerierten Influencer, auch virtuelle Influencer genannt (Grand View Research, 2024). Mit der Verbreitung neuer KI-Technologien, wie Sprachmodellen und Bildgeneratoren, ist der Übergang von virtuellen Influencern zu KI-Influencern, also virtuellen Influencern, die KI-Technologien zur Content-Erstellung nutzen, fließend. Ein prominentes Beispiel für KI-Influencer ist Lil Miquela, die auf Instagram 2,5 Millionen Follower zählt und mit Marken wie BMW und Givenchy kooperiert, sie wird von einer Agentur aus Los Angeles gesteuert. Sie beschreibt sich als „21-year-old Robot living in LA“ (Miquela, 2024). Ein weiteres Beispiel ist Alice Shaw mit 1 Million Followern, die Inhalte zu Fitness- und Mode-Themen postet und sich als „virtual soul powered by AI“ bezeichnet (Shaw, 2024).

Das Forschungsfeld der KI-Influencer ist noch wenig entwickelt. Erste Studien zeigen jedoch bereits, dass KI-Influencer sowohl positive als auch negative Werbeeffekte erzielen können (z. B. Franke et al., 2023; Sands et al., 2022; Thomas & Fowler, 2021; Lou et al., 2023). Die Autoren betonen die Notwendigkeit, weitere relevante Kennzahlen zur Werbeeffektivität zu untersuchen und weitere Produktkategorien einzubeziehen, um tiefergehende Erkenntnisse zu gewinnen und Unternehmen in ihrer Entscheidungsfindung besser zu unterstützen. Aus diesen Forschungslücken ergibt sich die folgende Forschungsfrage: *Wie beeinflusst der Einsatz von KI-Influencern im Vergleich zu menschlichen Influencern die Werbewirksamkeit von Instagram-Werbeanzeigen und welche Rolle spielt dabei die Produktkategorie?*

Ziel dieser Arbeit ist es, zu verstehen, wie sich KI-Influencer in ihrer Werbewirkung von menschlichen Influencern unterscheiden und welchen Einfluss die Produktkategorie darauf hat. Die Studie erweitert den bisherigen Forschungsstand, indem sie die Auswirkungen des Influencer-Typs auf Marken-Folgeabsicht,

Engagement-Rate und parasoziale Beziehungen (einseitige emotionale Bindung) untersucht. Zusätzlich werden zwei Produkte einbezogen, um die Auswirkungen auf Kongruenz (Fit zwischen Influencer und Produkt), Marken-Einstellung und Anzeigen-Einstellung in neuen Kontexten zu prüfen. Der Wirkungsmechanismus auf die Anzeigen-Einstellung wird dabei anhand der bisher noch nicht bei KI-Influencern quantitativ untersuchten parasozialen Beziehungen und der bereits erforschten Kongruenz beleuchtet. Zur Beantwortung der Forschungsfrage wird der aktuelle Forschungsstand zur Werbewirksamkeit dargestellt und die daraus abgeleiteten Hypothesen werden mithilfe eines 2 (Influencer-Typ: KI, Mensch) x 2 (Produktkategorie: Technik: Spotify, Kosmetik: Labello) Between-Subjects-Designs überprüft. Diese Arbeit leistet einen Beitrag zu einem tieferen Verständnis der Wirkung und Wirkmechanismen von KI-Influencern, erweitert die bestehende Forschung um relevante Kenngrößen und unterstützt ein besseres Verständnis der Interaktionen zwischen Mensch und KI.

Die Arbeit ist wie folgt aufgebaut: In der *Einleitung* wird das Forschungsthema vorgestellt, die Relevanz der Untersuchung aufgezeigt und die Forschungsfrage formuliert. Im Kapitel *Theoretische Grundlagen* wird der aktuelle Forschungsstand zu KI-Influencern analysiert, Hypothesen werden abgeleitet und in einem Forschungsmodell dargestellt. Die *Methodik* beschreibt den Aufbau und die Durchführung des 2 × 2-Experiments, einschließlich Operationalisierungen und Stichprobenbeschreibung. Im Kapitel *Ergebnisse* werden die Daten deskriptiv und mittels Hypothesentests ausgewertet. Im letzten Kapitel *Diskussion* werden die Resultate zusammengefasst und interpretiert, die praktischen und wissenschaftlichen Implikationen abgeleitet und Limitationen aufgezeigt, sowie Vorschläge für die zukünftige Forschung gemacht.

2. Theoretischer Hintergrund und Hypothesenbildung

Zunächst werden die zentralen Begriffe definiert, anschließend folgt eine tabellarische Zusammenstellung der wesentlichen Merkmale der analysierten Studien (siehe Tabelle 1). Im Anschluss werden die untersuchten zentralen Variablen gruppiert und strukturiert synthetisiert, um den aktuellen Forschungsstand darzustellen. Abschließend werden die Hypothesen eingeführt und im Forschungsmodell dargestellt.

2.1 Definition

Nach Kumar et al. (2019) wird KI wie folgt definiert: „AI refers to the broad idea that computers, through the use of software and algorithms, can think and perform tasks like humans“ (S. 135). Diese Technologie machen sich KI-Influencer zunutze, um mithilfe von Software menschenähnlichen Content zu generieren. Thomas und Fowler (2021) definieren KI-Influencer als „digitally created artificial humans, associated with Internet fame, using software and algorithms to perform tasks like humans“ (S. 12). Sands et al. (2022) unterscheiden weiter zwischen autonomen KI-Influencern, die vollständig von KI gesteuert werden und nicht-autonomen KI-Influencern, die unter menschlicher oder agenturbasierter Kontrolle stehen. Heutige KI-Influencer agieren nicht autonom und sind auf menschliche Steuerung angewiesen (Thomas & Fowler, 2021). Dies ist vermutlich darauf zurückzuführen, dass KIs selbst keine Rechtspersönlichkeit besitzen (Europäische Union, 2024) und daher weder eigenständig Werbeverträge abschließen noch Verantwortung für die erstellten Inhalte übernehmen können. Zudem besitzen KIs noch kein eigenes Bewusstsein, können den menschlichen Verstand nicht vollständig nachahmen und verfügen nicht über menschenähnliche Entscheidungsfähigkeiten (Google Cloud, 2024).

In der Literatur wird häufig der Begriff *virtueller Influencer* verwendet, obwohl den Probanden in Studien tatsächlich KI-Influencer wie Lil Miquela präsentiert oder ihnen ausdrücklich als solche vorgestellt werden (z. B. Sands et al., 2022; Franke et al., 2023; Belanche et al., 2024). So zeigten Song et al. (2024, S. 7) den Probanden die Beschreibung „Lil Miquela is one of the metahumans created by AI“, verwendeten jedoch weiterhin den allgemeinen Begriff virtueller Influencer. In der Studie von Kim et al. (2024) erhielten die Probanden Hinweise wie „she is an AI influencer“ (S. 8), die klar auf den KI-Hintergrund hinwiesen, während die Autoren beim Begriff virtueller Influencer blieben. Auch Volles et al. (2024) verweist durch Formulierungen wie „powered by artificial intelligence (AI)“ (S. 10) auf den Einsatz von KI-Technologien.

Um diese begriffliche Unschärfe aufzulösen und die spezifischen Eigenschaften von KI-Influencern klarer untersuchen zu können, wird in dieser Arbeit der Begriff *KI-Influencer* verwendet und den Probanden auch als solcher vermittelt. Dieser beschreibt

virtuelle, nichtautonome Influencer, die menschenähnlich auftreten und mithilfe von KI-Technologien Inhalte wie Fotos, Videos und Texte für ihre Follower und Werbepartner erstellen. Diese Definition bildet gleichzeitig die Grundlage für die unabhängige Variable Influencer-Typ.

2.2 Studienüberblick

Zur Kategorisierung und übersichtlichen Darstellung des aktuellen Forschungsstands zu KI-Influencern wurde eine Literaturtabelle erstellt (siehe Tabelle 1), die als Basis für weiterführende Analysen dient. Im Folgenden werden die Spalten beschrieben und besondere Merkmale hervorgehoben:

- Quelle: Insgesamt wurden 13 Studien gefunden und analysiert. Davon stammen 9 aus 2024, 2 aus 2023, 1 aus 2022 und 1 aus 2021.
- Land: Damit ist der Erhebungsort gemeint, also das Land, in dem die Teilnehmer befragt wurden: USA, Deutschland, UK, China, Singapur, Italien.
- Methode: 10 Studien führten Laborexperimente durch, während 3 Studien auf Interviews setzten.
- Unabhängige Variablen: Die unabhängigen Variablen umfassen den Influencer-Typ (virtuell, KI, menschlich, cartoonhaft), die Produktkategorie, den Botschaftsfokus, die Steuerung und Transgressionen.
- Abhängige Variablen: Insgesamt wurden 20 abhängige Variablen untersucht, wobei die Kaufabsicht, die Influencer-Folgeabsicht und die empfundene Unheimlichkeit am häufigsten betrachtet wurden.
- Produktkategorie: Mode, Beauty und Technik.
- Werbemittel: Die Stimuli (Instagram-Posting, Inserat und Twitter-Posting), über die die Werbebotschaft des KI-Influencers präsentiert wurde.
- Influencer: Dieser Punkt beschreibt den KI-Influencer, den die Probanden beurteilen sollten. In 10 Studien wurde Lil Miquela untersucht, wobei der Name in drei Fällen verfälscht wurde.
- Kernaussagen: Hier werden die Ergebnisse der Studien zusammengefasst.

Tabelle 1: Überblick über den aktuellen Forschungsstand zu KI-Influencern

Quelle	Land	Methode	Unabhängige Variablen	Abhängige Variablen	Produkt-kategorie	Werbe-mittel	Influencer	Kernaussagen
Thomas & Fowler (2021); <i>Journal of Advertising</i>	USA	Labor-experiment	Endorser-Typ (KI vs. Mensch), Transgression (vorhanden vs. nicht vorhanden)	Marken-Einstellung, Kaufabsicht	Sonnenbrille (Fiktive Marke)	Instagram-Posting	Fiktiver KI-Influencer	KI-Influencer sowie Prominente senken die Marken-Einstellung und Kaufabsicht, wenn sie eine Verfehlung begehen. Eine Verbesserung der Wahrnehmung nach der Verfehlung gelingt nur, wenn der KI-Influencer durch einen Prominenten ersetzt wird.
Sands et al. (2022); <i>European Journal of Marketing</i>	USA	Labor-experiment	Influencer-Typ (KI vs. Mensch), Influencer-Steuerung (Agentur vs. autonom)	Quellenvertrauen, WOM, Influencer-Folgeabsicht Sozialpsychologische Distanz, Personalisierung, Kommerzialisierung	Mode (Fiktive Marke)	Instagram-Posting	Lil Miquela	KI-Influencern wird weniger hinsichtlich ihrer Empfehlungen vertraut und die soziale Distanz ist höher, dennoch erzeugen sie mehr Word-of-Mouth und eine ähnlich hohe Bereitschaft, ihnen zu folgen.
Franke et al. (2023); <i>Journal of Advertising</i>	Deutschland	Labor-experiment	Influencer-Typ (Virtuell vs. Mensch), Produktkategorie (Technik vs. Kosmetik)	Anzeigen/Influencer-Einstellung, Kongruenz, Innovation Unheimlichkeit, Kaufabsicht	Lotion (Calvin Klein), Smart Speaker (Samsung)	Inserat mit Instagram-Influencer	Lil Miquela (Fiktiver Name)	Virtuelle Influencer erzielen geringere positive Einstellungen als menschliche, steigern jedoch die empfundene Neuartigkeit und erhöhen die Markeninnovation.

Tabelle 1: Fortsetzung

Quelle	Land	Methode	Unabhängige Variablen	Abhängige Variablen	Produkt-kategorie	Werbe-mittel	Influencer	Kernaussagen
Lou et al. (2023); <i>Journal of Advertising</i>	Singapur	Interview	-	-	-	Instagram-Posting	Lil Miquela	Virtuellen Influencern wird aufgrund von Neuheit, Unterhaltung, Information, Ästhetik und sozialer Interaktion gefolgt. Sie steigern die Markenbekanntheit und das Markenimage, erzielen jedoch aufgrund geringer Authentizität und schwacher parasozialer Beziehungen niedrigere Kaufabsichten.
Belanche et al. (2024); <i>Journal of Business Research</i>	USA	Labor-experiment	Influencer-Typ (Virtuell vs. Mensch), Produktkategorie (hedonisch vs. utilitaristisch)	Identifikation, Nützlichkeit, Quellenvertrauen	Laptop, Hotelzimmer (Fiktive Marken)	Instagram-Posting	Lil Miquela	Virtuelle Influencer werden als nützlicher wahrgenommen, dies gilt speziell bei utilitaristischen Produkten, während Follower eine stärkere Identifikation mit menschlichen Influencern zeigen.
Franke & Groeppel-Klein (2024); <i>Journal of Business Research</i>	Deutschland	Labor-experiment	Virtueller-Influencer-Typ (menschlich vs. cartoonhaft), Botschaftsfokus (Nutzen vs. Umsetzung)	Anzeigen-Einstellung, Sozialpsychologische Distanz, Glaubwürdigkeit, Neuheit	Parfüm (Calvin Klein), Smartwatch (Samsung)	Inserat mit Instagram-Influencer	Lil Miquela (Fiktiver Name)	Cartoonartige virtuelle Influencer werden als neuartiger empfunden, während menschlichen virtuellen Influencern mehr vertraut wird. Eine höhere psychologische Distanz senkt das Vertrauen, steigert jedoch die empfundene Neuartigkeit.

Tabelle 1: Fortsetzung

Quelle	Land	Methode	Unabhängige Variablen	Abhängige Variablen	Produkt-kategorie	Werbe-mittel	Influencer	Kernaussagen
Gutuleac et al. (2024); <i>Psychology & Marketing</i>	UK	Labor-experiment	Virtueller-Influencer-Typ (menschlich vs. cartoonhaft), Anthro-pomorphismus (hoch vs. niedrig)	Unheimlichkeit, Kaufabsicht	Kappen (Fiktive Marke)	Instagram-Posting	Lil Miquela, Imma.Gram, Nobody Sausage	Virtuelle Influencer mit stark menschenähnlichen Merkmalen werden als unheimlicher empfunden, was zu einer niedrigeren Kaufabsicht führt. Soziale Signale, wie Interaktionen mit anderen Personen, schwächen diese negative Wirkung ab.
Kim et al. (2024); <i>Journal of Business Research</i>	USA	Labor-experiment	Virtueller-Influencer-Typ (menschlich vs. cartoonhaft), Menschliches Verhalten (gering vs. ausgeprägt)	Unheimlichkeit, Coolness, Influencer-Folgeabsicht, Kaufabsicht	Mode, Kosmetik (Fiktive Marken)	Instagram-Posting	Lil Miquela, Noonouri (Fiktive Namen)	Virtuelle Influencer mit realistischem Aussehen und Verhalten werden als unheimlicher und cooler wahrgenommen. Unheimlichkeit reduziert die Folge- und Kaufabsicht, während Coolness diese positiv beeinflusst.
Koles et al. (2024); <i>Journal of Business Research</i>	Italien, Weltweit (Experten)	Interview	-	-	Mode, Kosmetik	-	Lil Miquela	Virtuelle Influencer haben eine True-to-Ideal (TTI) Authentizität, es fehlt ihnen jedoch an True-to-Fact (TTF) und True-to-Self (TTS) Authentizität, dies senkt ihre Glaubwürdigkeit und Effektivität.

Tabelle 1: Fortsetzung

Quelle	Land	Methode	Unabhängige Variablen	Abhängige Variablen	Produkt-kategorie	Werbe-mittel	Influencer	Kernaussagen
Mrad et al. (2024); <i>Journal of Advertising</i>	Weltweit	Interview	-	-	Mode	Instagram-Posting	Follower von Lil Miquela	Virtuelle Influencer können Neid, Bewunderung, Motivation und Dankbarkeit auslösen und parasoziale Beziehungen und soziale Vergleiche erzeugen.
Song et al. (2024); <i>Journal of Business Research</i>	China, USA	Labor-experiment, Textanalyse,	Influencer-Typ (Virtuell vs. Mensch)	Authentizität, Marken-Authentizität	Notizbücher, Smartwatch (Swatch), Sonnenbrille (Gucci), Müsli	Instagram-Posting, Inserat	Lil Miquela	Virtuelle Influencer werden als weniger authentisch wahrgenommen, dies beeinflusst die Markenauthentizität negativ. Ästhetische Markel oder Mehrfach-Sponsoring können diesen Effekt abschwächen.
Volles et al. (2024); <i>Journal of Business Research</i>	USA	Labor-experiment	Influencer-Typ (Virtuell vs. Mensch)	Kaufabsicht, Influencer-Folgeabsicht	Handtaschen (Louis Vuitton, Gucci), Schokoriegel (Fiktive Marke)	Instagram-Posting	Fiktiver virtueller Influencer	Virtuelle Influencer erhöhen die wahrgenommene Einzigartigkeit, dies steigert die Kauf- und Folgebereitschaft, mangelndes Vertrauen schwächt diese Wirkung wiederum ab.
Zhou et al. (2024); <i>Journal of Marketing</i>	-	Labor-experiment, Feld-experiment	Influencer-Typ (Virtuell vs. Mensch), Sensorische Erfahrung (proximal vs. distal)	Kaufabsicht, Produkt-Einstellung	Hotel, T-Shirt, Schlafkopfhörer (Fiktive Marken)	Twitter-Posting	Fiktiver virtueller Influencer	Virtuelle Influencer erzielen bei Produkten mit proximalen sensorischen Eigenschaften, wie Berührung und Geruch, geringere Kaufabsichten, da ihre sensorischen Fähigkeiten als geringer wahrgenommen werden.

Die untersuchten Effekte werden in die thematischen Gruppen *Wahrnehmung* (z. B. Unheimlichkeit, Coolness, Neuheit), *psychologische Mechanismen* (z. B. Authentizität, sozialpsychologische Distanz, parasoziale Beziehungen), *Einstellungen* (Marken-Einstellung und Anzeigen-Einstellung) und *Kauf- und Verhaltensabsichten* (Kaufabsicht und Influencer-Folgeabsicht) unterteilt. In den ersten beiden Kapiteln werden die Wahrnehmung von KI-Influencern und die zugrunde liegenden psychologischen Mechanismen analysiert. In den Kapiteln zu *Einstellungen* sowie *Kauf- und Verhaltensabsichten* wird die Werbewirksamkeit untersucht.

2.3 Wahrnehmung

2.3.1 Erkennung von KI-Influencern



Abbildung 1: Instagram-Post (Miquela, 2023) Abbildung 2: Instagram-Post (Lopez, 2024)

Menschen erkennen KI-Influencerinnen wie Lil Miquela nicht immer als solche, wenn keine explizite KI-Kennzeichnung vorliegt. In Abbildung 1 ist Lil Miquela zu sehen, deren künstliche Merkmale wie die symmetrischen Sommersprossen und die gezeichnet wirkenden Augenbrauen ihre virtuelle Natur andeuten. Im Gegensatz dazu wäre die KI-Influencerin Aitana Lopez (Abbildung 2) kaum von einem echten Menschen zu unterscheiden. Franke et al. (2023) untersuchten diese Thematik, indem sie Probanden Lil Miquela präsentierten und fragten, ob sie die Influencerin, als menschlich oder KI einordnen. Die Ergebnisse zeigten, dass 32 % Lil Miquela fälschlicherweise für menschlich hielten, während 29 % sich unsicher waren. Lou et al. (2023) interviewten Follower von Lil Miquela, die deren virtuelle Natur zwar meist

erkannten, jedoch aufgrund ihrer menschenähnlichen Gestaltung gelegentlich daran erinnert werden mussten. Eine Teilnehmerin beschrieb dies treffend: „[Lil Miquela] she’s very close [to humans], but still not there“ (Lou et al., 2023). Realistischere Darstellungen wie die von Aitana Lopez wären hingegen ohne vorherige Kenntnis kaum als KI-Influencerinnen erkennbar, was zeigt, dass die Erkennung maßgeblich von der Gestaltung abhängt. Unsicherheit über die Natur des Influencers kann dabei negative Folgen wie ein verstärktes Empfinden von Unheimlichkeit hervorrufen und die Wahrnehmung insgesamt beeinflussen (Franke et al., 2023).

2.3.2 Unheimlichkeit

Die empfundene Unheimlichkeit von KI-Influencern, die menschenähnlich sind, wirkt sich negativ auf die Effektivität von Werbung aus (Franke et al., 2023; Lou et al., 2024; Gutuleac et al., 2024; Kim et al., 2024) und kann auf die Uncanny-Valley-Theorie von Mori et al. (2012) zurückgeführt werden. Diese beschreibt die Reaktion auf nahezu menschenähnliche virtuelle Abbildungen, bei der die Sympathie zunächst mit zunehmender Menschlichkeit steigt. Bis der kritische Punkt erreicht wird, an dem die Akzeptanz plötzlich in Ablehnung umschlägt, danach steigt die Akzeptanz mit zunehmender Menschlichkeit wieder an, was als das Tal der Unheimlichkeit bekannt ist. Dieser Effekt tritt auf, wenn die Darstellung fast vollständig menschenähnlich ist, aber noch kleine künstliche Details erkennbar bleiben oder Unsicherheit über die Menschlichkeit besteht.

Kim et al. (2024) wiesen nach, dass KI-Influencer mit niedrigem Formalrealismus als wenig unheimlich empfunden wurden. Bei mittlerem Formrealismus, also nahezu menschenähnlich, aber mit erkennbaren künstlichen Merkmalen, zu dem auch Lil Miquela zählt, stieg die empfundene Unheimlichkeit an. Hochrealistische KI-Influencer wurden hingegen wieder als weniger unheimlich wahrgenommen. Dieses Muster verdeutlicht die Uncanny-Valley-Theorie und zeigt, dass Lil Miquela in diesen kritischen Bereich fällt.

Franke et al. (2023) untersuchten dieses Phänomen in Bezug auf die Unsicherheit bei der Bestimmung des Influencer-Typs anhand von Lil Miquela. Dabei stellten sie fest, dass die empfundene Unheimlichkeit abnahm, wenn die Influencerin als KI-Influencerin gekennzeichnet war, wodurch Unsicherheiten bei den Probanden

reduziert wurden. Lou et al. (2023) interviewten Follower von Lil Miquela und fanden heraus, dass die Unheimlichkeit durch das gezielte Hinzufügen visueller Makel, wie Falten, gemindert werden kann, da dies die Menschenähnlichkeit steigert.

Gutuleac et al. (2024) bewiesen, dass die empfundene Unheimlichkeit abnimmt, wenn KI-Influencer, in diesem Fall Lil Miquela, in Fotos zusammen mit echten Menschen dargestellt werden. Die Autoren erklärten diesen Effekt mit der CASA-Theorie (Computers as Social Actors) von Nass und Moon (2000). Diese besagt, dass digitale Personen als menschlicher angesehen werden, wenn sie soziale Signale zeigen, zum Beispiel wenn sie zusammen mit realen Menschen abgebildet sind. Eine weitere Erklärung gaben Gutuleac et al. (2024) mit der Schema-Theorie von Nass und Moon (2000). Demnach nutzen Menschen ihr bestehendes Wissen und ihre Schemata, um Neues einzuordnen. Wird ein KI-Influencer neben realen Personen dargestellt, können die erworbenen Muster auf den KI-Influencer übertragen werden, was zu einer Steigerung seiner menschlichen Anmutung führt. Eine Verringerung der Unheimlichkeit beobachteten auch Kim et al. (2024), wenn sich der KI-Influencer menschenähnlicher verhielt, indem er personalisierter und ausführlicher auf Kommentare reagierte.

2.3.3 Coolness, Neuheit und Innovation

KI-Influencer werden als cooler wahrgenommen, wenn sie einen höheren Formrealismus aufweisen und personalisierte sowie ausführliche Antworten auf Kommentare geben (Kim et al., 2024). Anzeigen von KI-Influencern werden zudem als neuartiger empfunden, was die wahrgenommene Innovationskraft der Marke steigert (Franke et al., 2023). Lou et al. (2023) fanden heraus, dass Neuheit neben Information, Unterhaltung, Neugier, Ästhetik und sozialer Interaktion einer der Hauptgründe ist, warum KI-Influencern gefolgt wird.

2.3.4 Personalisierung und Kommerzialisierung

In Bezug auf die Fähigkeit zur Personalisierung von Inhalten stellten Sands et al. (2022) fest, dass KI-Influencer ähnlich wie menschliche Influencer wahrgenommen werden. Die empfundene Kommerzialisierung ist bei KI-Influencern stärker, insbesondere wenn sie von Agenturen gesteuert werden, im Vergleich zu einer autonomen Steuerung.

2.3.5 Nützlichkeit

Belanche et al. (2024) prüften die wahrgenommene Nützlichkeit von Instagram-Postings, die sowohl hedonische Produkte (z. B. Hotelbuchungen) als auch utilitaristische Produkte (z. B. Laptops) bewarben. Ihre Ergebnisse zeigten, dass Postings von KI-Influencern zu utilitaristischen Produkten als nützlicher wahrgenommen wurden als Postings zu hedonischen Produkten. Zudem ergab die Studie, dass Postings zu utilitaristischen Produkten von KI-Influencern als nützlicher eingeschätzt wurden als entsprechende Postings von menschlichen Influencern. Laut Belanche et al. (2024) lässt sich dies darauf zurückführen, dass Konsumenten bei utilitaristischen Produkten rationale Informationen bevorzugen (Holbrook & Hirschman, 1982). Zudem wird Postings von KI-Influencern aufgrund ihrer systematischen, analytischen Ansätze sowie ihres Zugriffs auf umfangreiche Daten eine höhere Kompetenz zugeschrieben (Huang & Rust, 2018).

Fazit: Die Ergebnisse zur Wahrnehmung zeigen, dass soziale Signale, wie die Interaktion mit realen Personen (Gutuleac et al., 2024), die Personalisierung der Kommunikation (Kim et al., 2024), die klare Kennzeichnung als KI (Franke et al., 2023) und visuelle Anpassungen wie das Hinzufügen von Makeln (Lou et al., 2023) die Unheimlichkeit von KI-Influencern verringern und somit ihre Wahrnehmung und Werbewirksamkeit verbessern können. KI-Influencer werden als cooler wahrgenommen (Kim et al., 2024) und ihre Anzeigen gelten als neuartiger (Franke et al., 2023; Lou et al., 2023) und nützlicher (Belanche et al., 2024), insbesondere wenn mit utilitaristischen Produkten geworben wird. Gleichzeitig wird deutlich, dass Menschen KI-Influencer nicht immer als solche identifizieren (Franke et al., 2023; Lou et al., 2023).

2.4 Psychologische Mechanismen

2.4.1 Sozialpsychologische Distanz

Die sozialpsychologische Distanz beschreibt, wie weit entfernt andere im Vergleich zum eigenen Selbst wahrgenommen werden, wobei größere Distanzen eine abstraktere Wahrnehmung und weniger persönliche Verbindung bedeuten (Trobe & Liberman, 2010). Sands et al. (2022) haben gezeigt, dass KI-Influencer im Vergleich zu

menschlichen Influencern eine höhere sozialpsychologische Distanz aufweisen. Dies ist auf ihre fehlende physische Präsenz und geringere Nahbarkeit zurückzuführen. Franke und Groeppel-Klein (2024) fanden zudem heraus, dass menschenähnliche KI-Influencer im Vergleich zu cartoonartigen Influencern eine geringere sozialpsychologische Distanz aufwiesen.

2.4.2 Authentizität

Lou et al. (2023) fanden mittels Interviews heraus, dass KI-Influencer im Vergleich zu menschlichen Influencern als weniger authentisch wahrgenommen werden. Als Gründe identifizierten sie die künstliche Natur der KI-Influencer und deren mangelnde Fähigkeit, Produkte real zu testen oder persönliche Erfahrungen zu teilen. Koles et al. (2024) bestätigen dies ebenfalls mit Interviews und betrachten die Authentizität dreigeteilt. Sie zeigen, dass KI-Influencer eine Authentizität, die dem Ideal entspricht (true to ideal), erreichen können, indem sie den Erwartungen entsprechen. Allerdings fehlt es ihnen an Authentizität, die den Fakten entspricht (true to fact), da ihre Darstellungen nicht auf echten Erfahrungen beruhen. Zudem besitzen sie keine Authentizität, die dem Selbst entspricht (true to self), da sie keine intrinsischen Motive haben. Song et al. (2024) haben diesen Effekt empirisch untersucht und bestätigten die geringere Authentizität von KI-Influencern. Sie zeigten zudem, dass eine reduzierte Authentizität negativ auf die Marken-Authentizität wirkt. Dieser negative Einfluss konnte jedoch gemildert werden, wenn dem KI-Influencer optische Merkmale wie Sommersprossen hinzugefügt wurden oder wenn der KI-Influencer für mehrere Marken gleichzeitig warb.

2.4.3 Vertrauen

Das Quellenvertrauen, also das Vertrauen in die Produktempfehlungen, ist bei KI-Influencern im Vergleich zu menschlichen Influencern geringer und wird durch die sozialpsychologische Distanz vermittelt (Sands et al., 2022). Belanche et al. (2024) untersuchten ebenfalls das Vertrauen in die Produktempfehlungen und kamen zu dem Schluss, dass die wahrgenommene Nützlichkeit und die Identifikation der Nutzer positiv auf das Vertrauen wirken.

2.4.4 Identifikation

Belanche et al. (2024) wiesen nach, dass sich Probanden stärker mit menschlichen Influencern identifizieren konnten als mit KI-Influencern, unabhängig davon, ob utilitaristische oder hedonistische Produkte beworben wurden. Diesen Effekt erklären Belanche et al. (2024) mit der Social Identity Theory von Tajfel und Turner (1986), die besagt, dass Menschen sich zu Gruppen hingezogen fühlen, denen sie selbst ähnlich sind. In diesem Fall empfinden sie menschliche Influencer ähnlicher zu sich selbst, was die Identifikation erhöht. Lou et al. (2023) kamen in ihren Interviews zu dem Ergebnis, dass sich die Probanden weniger mit KI-Influencern identifizieren, da sich deren Lifestyle von dem der Probanden unterscheidet.

2.4.5 Parasoziale Beziehungen

Parasoziale Beziehungen sind einseitige, emotionale Bindungen zu Personen, die man nicht persönlich kennt und mit denen keine reale Interaktion stattfindet (Horton & Wohl, 1956). Lou et al. (2023) und Koles et al. (2024) untersuchten diese parasozialen Beziehungen zu KI-Influencern und fanden in ihren Interviews heraus, dass diese Beziehungen schwächer sind als zu menschlichen Influencern. Begründet wurde dies damit, dass KI-Influencer als weniger authentisch wahrgenommen werden und keine echten Emotionen oder persönlichen Erfahrungen zeigen (Lou et al., 2023). Im Gegensatz dazu beobachteten Mrad et al. (2023), dass KI-Influencer Emotionen wie Sympathie bei ihren Followern hervorrufen und parasoziale Beziehungen aufbauen können, auch wenn diese aufgrund fehlender Authentizität und fehlender echter Erlebnisse weniger ausgeprägt sind.

Fazit: KI-Influencer werden als weniger authentisch wahrgenommen (Sands et al., 2022; Lou et al., 2023; Koles et al., 2024; Song et al., 2024) und weisen eine höhere sozialpsychologische Distanz auf (Sands et al., 2022). Follower vertrauen den Produktempfehlungen von KI-Influencern weniger (Sands et al., 2022; Belanche et al., 2024) und identifizieren sich weniger mit ihnen (Sands et al., 2022; Lou et al., 2023). Folglich bauen Follower schwächere parasoziale Bindungen zu KI-Influencern auf (Lou et al., 2023; Koles et al., 2024). Dennoch sind KI-Influencer in der Lage, solche Beziehungen zu entwickeln, wenn auch nicht so ausgeprägt (Mrad et al., 2023).

2.5 Einstellungen

2.5.1 Marken-Einstellung

Thomas und Fowler (2021) zeigten, dass KI-Influencer die Marken-Einstellung ähnlich positiv beeinflussen können wie menschliche Influencer. Sie stellten fest, dass die Marken-Einstellung sank, wenn ein KI-Influencer eine Transgression beging, indem er eine normverletzende Aussage tätigte. Dieser negative Effekt zeigte sich auch bei menschlichen Influencern. Allerdings verstärkte der Austausch eines KI-Influencers durch einen anderen KI-Influencer den negativen Effekt, während der Wechsel zu einem menschlichen Influencer positive Auswirkungen auf die Marken-Einstellung hatte. Dies wurde darauf zurückgeführt, dass KI-Influencer als weniger unabhängig und als ersetzbar wahrgenommen werden, wodurch die Marke weiterhin für das Verhalten verantwortlich gemacht wurde.

Lou et al. (2023) stellten fest, dass KI-Influencer häufig mit Innovation und Neuheit in Verbindung gebracht werden, wodurch die beworbene Marke moderner und wettbewerbsfähiger erscheint. Franke et al. (2023) bestätigten in ihrer Untersuchung diese empfundene Neuheit. Lou et al. (2023) berichteten zudem, dass sich die Befragten vorstellen können durch KI-Influencer neue Marken zu entdecken.

2.5.2 Anzeigen-Einstellung

Franke et al. (2023) fanden keinen direkten Effekt des Influencer-Typs auf die Anzeigen-Einstellung, jedoch zeigen die Mittelwerte, dass menschliche Influencer höhere Werte erzielten. Die Influencer-Einstellung vermittelte dabei den positiven Effekt auf die Anzeigen-Einstellung vollständig. Zusätzlich hatte die empfundene Expertise einen positiven Einfluss auf die Anzeigen-Einstellung. Franke und Groeppel-Klein (2024) verglichen menschenähnliche KI-Influencer mit cartoonartigen KI-Influencern und fanden heraus, dass ein menschenähnlicher KI-Influencer die Vertrauenswürdigkeit erhöht, was die Anzeigen-Einstellung wiederum erhöht. Ein cartoonartiger KI-Influencer hingegen erzeugt eine höhere Neuheit, die sich ebenfalls positiv auf die Anzeigen-Einstellung auswirkt.

Fazit: KI-Influencer können, ähnlich wie menschliche Influencer, die Marken-Einstellung positiv beeinflussen (Thomas und Fowler, 2021), erzielen jedoch eine

geringere Anzeigen-Einstellung (Franke et al., 2023). Zudem können sie zur Steigerung der Markenbekanntheit und des Markenimages beitragen (Lou et al., 2023).

2.6 Kauf- und Verhaltensabsichten

2.6.1 Kaufabsicht

Thomas und Fowler (2021) zeigten, dass Transgressionen die Kaufabsicht bei KI-Influencern reduzieren, wobei dieser Effekt nur durch den Wechsel zu einem menschlichen Influencer gemildert werden kann. Franke et al. (2023) fanden heraus, dass eine positive Anzeigen-Einstellung die Kaufbereitschaft steigert. Lou et al. (2023) identifizierten mit ihrem Interview mangelnde Authentizität und schwächere parasoziale Beziehungen als wesentliche Faktoren für eine geringere Kaufbereitschaft bei KI-Influencern. Volles et al. (2024) wiesen nach, dass KI-Influencer die wahrgenommene Einzigartigkeit steigern können, was die Kaufbereitschaft besonders bei starker Identifikation erhöht, jedoch durch fehlendes Vertrauen abgeschwächt wird. Zhou et al. (2024) stellten fest, dass die Kaufbereitschaft sinkt, wenn Produkte beworben werden, die proximale Sinne (z. B. Tasten, Riechen, Schmecken) erfordern, da KI-Influencer als weniger fähig wahrgenommen werden, diese Erfahrungen authentisch zu vermitteln.

2.6.2 Influencer-Folgeabsicht

Zwischen menschlichen und KI-Influencern traten bei der Influencer-Folgeabsicht keine Unterschiede auf, sodass beiden Influencern ähnlich gerne gefolgt wird. Dies könnte auf die empfundene Neuheit und die geringere soziale Distanz von KI-Influencern zurückzuführen sein (Sands et al., 2022). KI-Influencer erzielten zudem eine höhere Word-of-Mouth-Intention, die Sands et al. (2022) ebenfalls mit der empfundenen Neuheit argumentierten. Lou et al. (2023) bestätigten in ihren Interviews die Neuheit als Hauptgrund, um KI-Influencern zu folgen und identifizierten fünf weitere Motive: Unterhaltung, Informationssuche, Aktualität, ästhetische Inspiration und Gemeinschaftsgefühl.

Fazit: Transgressionen (Thomas & Fowler, 2021) und mangelnde Authentizität (Lou et al., 2023) reduzieren die Kaufbereitschaft bei KI-Influencern. Volles et al. (2024) betonten jedoch, dass wahrgenommene Einzigartigkeit und starke Identifikation die

Kaufbereitschaft erhöhen können. KI-Influencern wird genauso gerne gefolgt wie menschlichen Influencern und sie erzeugen mehr WOM (Sands et al., 2022).

2.7 Hypothesen

Aufbauend auf den theoretischen Grundlagen werden in diesem Abschnitt spezifische Hypothesen formuliert, die sich direkt auf die Forschungsfrage beziehen.

Thomas und Fowler (2021) identifizierten keine statistisch signifikanten Unterschiede zwischen KI- und menschlichen Influencern hinsichtlich der Marken-Einstellung. Ebenso fanden Franke et al. (2023) keinen direkten statistisch signifikanten Effekt des Influencer-Typs auf die Anzeigen-Einstellung, allerdings wurde die Influencer-Einstellung bei menschlichen Influencern besser bewertet, was wiederum einen positiven Effekt auf die Anzeigen-Einstellung hatte. Dies deutet darauf hin, dass menschliche Influencer bei der Anzeigen-Einstellung im Vorteil sein könnten, während KI-Influencer aufgrund ihrer Neuheit (Sands et al., 2022; Lou et al., 2023) und Coolness (Kim et al., 2024) möglicherweise bei der Engagement-Rate und der Marken-Folgeabsicht punkten könnten. Daraus ergibt sich folgende Hypothese:

***H1:** KI-Influencer erzielen im Vergleich zu menschlichen Influencern bei Werbeanzeigen höhere Werte hinsichtlich (a) der Marken-Folgeabsicht und (b) der Engagement-Rate, während sie bei (c) der Marken-Einstellung ähnliche Werte und bei (d) der Anzeigen-Einstellung niedrigere Werte erreichen.*

Kamins (1990) beschreibt den sogenannten Matchup-Effekt, wonach Werbeanzeigen dann besonders wirksam sind, wenn Endorser und beworbenes Produkt gut aufeinander abgestimmt sind. Dieser Effekt basiert auf der Annahme, dass Konsumenten positiver auf Werbung reagieren, wenn sie einen natürlichen Fit zwischen dem Endorser und dem Produkt wahrnehmen. Franke et al. (2023) zeigten, dass die Kongruenz zwischen KI-Influencern und technischen Produkten sowie zwischen menschlichen Influencern und kosmetischen Produkten statistisch signifikant höher ist als bei der jeweils anderen Produktkategorie. Dies deutet darauf hin, dass KI-Influencer aufgrund ihrer Assoziation mit Innovation (Franke et al., 2023) besonders gut mit technischen Produkten harmonieren. Ein weiteres Indiz liefert

Belanche et al. (2024), die feststellten, dass Werbeanzeigen für technische Produkte als nützlicher wahrgenommen werden. Im Gegensatz dazu werden bei kosmetischen Produkten, die häufig menschliche Attribute wie Körper und Haut betonen und proximale Sinne erfordern, menschliche Influencer bevorzugt, da KI-Influencern diese Fähigkeiten nicht zugeschrieben werden (Zhou et al., 2024). Diese Überlegungen und Ergebnisse führen zur folgenden Hypothese:

***H2:** Die Produktkategorie moderiert die Wirkung von Influencern auf die Ergebnisse von Werbeanzeigen hinsichtlich (a) der Marken-Einstellung, (b) der Anzeigen-Einstellung, (c) der Marken-Folgeabsicht und (d) der Engagement-Rate, wobei KI-Influencer bei technischen Produkten und menschliche Influencer bei kosmetischen Produkten bessere Ergebnisse erzielen.*

Franke et al. (2023) fanden heraus, dass keine signifikanten Unterschiede in der wahrgenommenen Kongruenz zwischen dem Influencer-Typ und den Produkten existieren, wenn nicht zwischen der technischen und kosmetischen Produktkategorie unterschieden wird. Folglich lautet die Hypothese:

***H3:** KI-Influencer und menschliche Influencer erzielen eine ähnliche Kongruenz, wenn die Produktkategorie nicht einbezogen wird.*

Wird zwischen der technischen und der kosmetischen Produktkategorie unterschieden, stellten Franke et al. (2023) signifikante Unterschiede in der wahrgenommenen Kongruenz fest. Dies deutet darauf hin, dass die Produktkategorie die Kongruenz beeinflusst bzw. moderiert. Aufgrund der Assoziation von KI-Influencern mit Innovation (Franke et al., 2023), ihrer empfundenen Neuheit (Sands et al., 2022; Lou et al., 2023) und ihrer wahrgenommenen Nützlichkeit (Belanche et al., 2024) passen sie besser zu technischen Produkten, was die Kongruenz erhöht und im Einklang mit dem Matchup-Effekt (Kamins, 1990) steht. Daher ergibt sich folgende Hypothese:

***H4:** Die Kongruenz zwischen Influencern und Produkt wird durch die Produktkategorie moderiert, sodass KI-Influencer eine höhere Kongruenz mit technischen Produkten erzielen und menschliche Influencer eine höhere Kongruenz mit kosmetischen Produkten erzielen.*

Die Kongruenz beeinflusst direkt als auch indirekt über die empfundene Expertise die Anzeigen-Einstellung (Franke et al., 2023). Die direkte Wirkung der Kongruenz auf die Anzeigen-Einstellung wird mit folgender Hypothese geprüft:

H5: Eine höhere Kongruenz führt zu einer positiveren Anzeigen-Einstellung.

Parasoziale Beziehungen sind bei KI-Influencern geringer ausgeprägt, dies wurde in mehreren Interviews festgestellt (Lou et al., 2023; Koles et al., 2024; Mrad et al., 2023). Die schwächere parasoziale Beziehung bei KI-Influencern könnte darauf zurückzuführen sein, dass sie ihren Followern weniger ähnlich sind und dadurch weniger Identifikationspotenzial bieten (Sands et al., 2022; Lou et al., 2023). Zudem werden sie oft als weniger authentisch wahrgenommen (Lou et al., 2023; Koles et al., 2024). Zudem werden KI-Influencer als weniger vertrauenswürdig wahrgenommen und sie weisen eine höhere sozialpsychologische Distanz zu ihren Followern auf (Sands et al., 2022). Daraus ergibt sich folgende Hypothese:

H6: Die Intensität der parasozialen Beziehung ist bei menschlichen Influencern höher als bei KI-Influencern.

Eine höhere parasoziale Beziehung hat positive Auswirkungen auf die Werbeeffektivität, da sie das Vertrauen stärkt (Chung & Cho, 2017). Außerdem steigert sie den Kundenwert (Yuan et al., 2019) und die Bereitschaft, Empfehlungen von Influencern anzunehmen (Conde & Casais, 2023). Zudem senkt sie die wahrgenommenen Eigennutzmotive (Aw & Chuah, 2021) und die Persuasionskenntnis, was zu höheren Kaufabsichten und Markenbewertungen führt (Breves et al., 2021). Daraus ergibt sich folgende Hypothese:

H7: Eine höhere parasoziale Beziehung führt zu einer positiveren Anzeigen-Einstellung.

Bei technischen Produkten könnten parasoziale Beziehungen einen geringeren Einfluss haben als bei kosmetischen Produkten. Technische Produkte erfordern häufig Informationen zu Funktionalität und Spezifikationen, die weniger durch emotionale

Bindungen geprägt sind. Nach dem Elaboration Likelihood Model (ELM) von Petty und Cacioppo (1986) verarbeiten Konsumenten solche technischen Informationen häufig über den zentralen Weg und beschäftigen sich intensiver mit den Details. Dies legt nahe, dass bei technischen Produkten, bei denen rationale Informationen bevorzugt werden (Holbrook & Hirschman, 1982), parasoziale Beziehungen möglicherweise an Bedeutung verlieren. Daher ergibt sich folgende Hypothese:

H8: Die Produktkategorie moderiert den Effekt der parasozialen Beziehung, indem sie bei technischen Produkten den Effekt der parasozialen Beziehung abschwächt.

Eine positivere Anzeigen-Einstellung erhöht die Marken-Einstellung (Lee und Koo, 2015; MacKenzie et al., 1986). Es kann angenommen werden, dass eine positivere Anzeigen-Einstellung auch die Marken-Folgeabsicht und die Engagement-Rate erhöht. Diese Überlegungen leiten zu der folgenden Hypothese:

H9: Eine positivere Anzeigen-Einstellung erhöht (a) die Marken-Einstellung, (b) die Marken-Folgeabsicht und (c) die Engagement-Rate.

2.8 Forschungsmodell

Die Hypothesen (H1–H9) sind im Forschungsmodell dargestellt (siehe Abbildung 3).

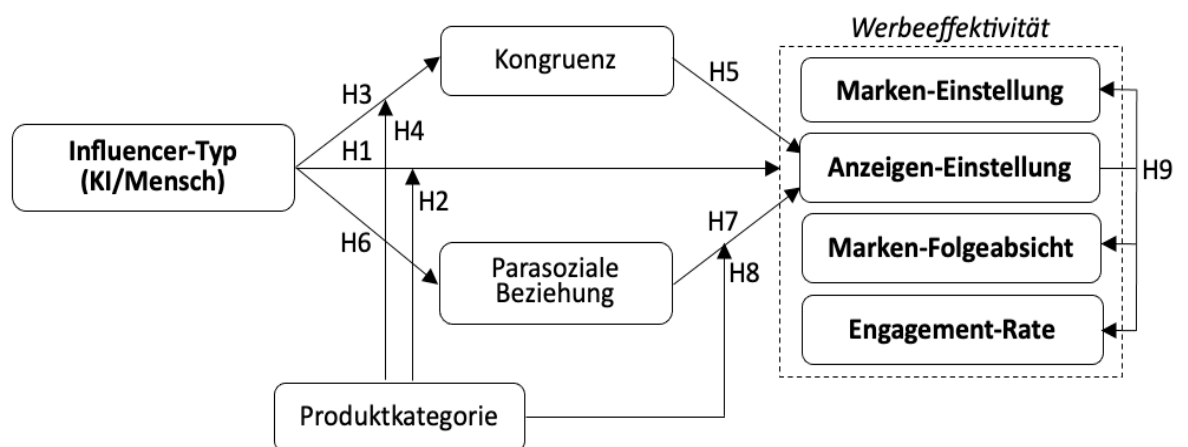


Abbildung 3: Forschungsmodell

3. Methodik

Zur Überprüfung der Hypothesen H1 bis H9 wurde ein 2×2 Between-Subjects-Design gewählt, das mithilfe einer Online-Befragung durchgeführt wurde. Im Kapitel *Stichprobenbeschreibung* wird die Rekrutierung und Zusammensetzung der Stichprobe beschrieben. Anschließend folgt das Kapitel *Operationalisierung*, in dem dargestellt wird, wie die abhängigen, mediierenden und kontrollierenden Variablen operationalisiert wurden. Im Kapitel *Experimentdesign* werden der Forschungsaufbau, die eingesetzten Stimuli (Influencer und Produktkategorie) und der genaue Ablauf des Experiments beschrieben.

3.1 Stichprobenbeschreibung

Die Rekrutierung der Probanden erfolgte über meine Social-Media-Profilen (Facebook, Instagram, LinkedIn), studentische WhatsApp- und Facebook-Gruppen sowie die Umfrageplattformen SurveySwap (SurveySwap, 2025) und SurveyCircle (SurveyCircle, 2025). Diese Plattformen arbeiten mit einem Punktesystem, bei dem Nutzer durch das Beantworten anderer Umfragen Punkte sammeln, die sie anschließend gegen qualifizierte Antworten auf ihre eigene Umfrage einlösen können. Zusätzlich wurde unter allen Probanden ein Gutschein nach Wahl verlost. Die Befragung richtete sich ausschließlich an Personen, die mit der Instagram-App vertraut waren, das heißt, sie hatten die Plattform bereits genutzt und waren mit deren Funktionen vertraut.

Von den insgesamt 300 erhaltenen Umfrageergebnissen waren 236 ($n = 236$) gültig. Aus der Analyse wurden die Ergebnisse von Probanden ausgeschlossen, die die Umfrage nicht vollständig ausgefüllt haben ($n = 25$), die die Mindestbearbeitungszeit von 90 Sekunden nicht erreicht oder die Maximalbearbeitungszeit von 10 Minuten überschritten haben ($n = 7$) oder die die nachfolgenden Kriterien nicht erfüllten. Dazu zählten Personen, die die Manipulationschecks nicht bestanden haben, also die den Beschreibungstext mit der KI-Information nicht aufmerksam gelesen hatten ($n = 2$), oder die Fragen nicht unter der Annahme, dass es sich um einen KI-Influencer handelt, beantwortet haben ($n = 17$). Ebenfalls ausgeschlossen wurden Ergebnisse von

Personen, die die präsentierte Marke nicht kannten ($n = 10$) oder angaben, den Influencer zu kennen ($n = 3$), um potenzielle Störeinflüsse zu vermeiden.

Die Zuweisung zu den vier experimentellen Gruppen (2×2) erfolgte zufällig. Um jedoch die zu erwartenden Ausschlüsse durch die Manipulationschecks in den KI-Influencer-Gruppen auszugleichen, wurde die Wahrscheinlichkeit, einer KI-Influencer-Gruppe zugeordnet zu werden, geringfügig erhöht. Diese Anpassung stellte ausgeglichene Gruppengrößen mit gültigen Antworten sicher. Die finalen Gruppen setzten sich wie folgt zusammen: KI-Technik ($n = 56$; KI-Influencer mit Spotify-Produkt), KI-Kosmetik ($n = 58$; KI-Influencer mit Labello-Produkt), Mensch-Technik ($n = 61$; menschlicher Influencer mit Spotify-Produkt) und Mensch-Kosmetik ($n = 61$; menschlicher Influencer mit Labello-Produkt).

Das Durchschnittsalter der Probanden betrug 27,93 Jahre ($SD = 5,23$), wobei 55,93 % weiblich, 43,22 % männlich und 0,85 % divers waren. Die Probanden stammten aus dem deutschsprachigen Raum, insbesondere aus Österreich und Deutschland. Die Bearbeitungsdauer betrug im Durchschnitt 2 Minuten und 51 Sekunden ($SD = 1$ Minute). Die Befragung wurde im Zeitraum vom 19. Jänner bis zum 31. Jänner 2025 durchgeführt. In diesem Zeitraum traten keine wesentlichen Ereignisse oder Skandale mit Bezug zu Influencern auf, die die Ergebnisse hätten beeinflussen können.

3.2 Operationalisierung

Ein Überblick über die Operationalisierung der abhängigen Variablen und Mediatoren findet sich in Tabelle 2. Die Kongruenz basiert auf einer 3-Item-7-Punkte-Likert-Skala nach Franke et al. (2023) und Fleck et al. (2012) (Cronbach's $\alpha = 0,85$). Die parasoziale Beziehung wurde durch eine 9-Item-7-Punkte-Likert-Skala nach Chung und Cho (2017) erfasst ($\alpha = 0,70$). Für die Marken-Folgeabsicht kam eine 1-Item-7-Punkte-Likert-Skala zum Einsatz, die auf Sands et al. (2022) basiert und dahingehend adaptiert wurde, dass die ursprüngliche Frage zur Influencer-Folgeabsicht auf die Marken-Folgeabsicht geändert wurde. Die Anzeigen-Einstellung stützt sich auf ein 4-Item-7-Punkte-Semantisches Differential nach Franke et al. (2023) und Silvera und Austad (2004) ($\alpha = 0,92$). Die Marken-Einstellung folgt einem 4-Item-7-Punkte-Semantisches Differential nach Thomas und Fowler (2021) und Till und Busler (2000) ($\alpha = 0,98$).

Die Engagement-Rate basiert auf einer selbst entwickelten 3-Item-Dichotomen-Skala mit den Fragen: Würden Sie diesen Instagram-Beitrag liken/kommentieren/teilen?

Die Kontrollvariablen wurden wie folgt erfasst: Das Product-Involvement (Franke et al., 2023) sowie die technische Affinität wurden jeweils mit einer 1-Item-7-Punkte-Likert-Skala gemessen. Die Influencer- und Markenbekanntheit wurden mit einer 1-Item-Dichotomen-Skala geprüft (war Ihnen die beworbene Marke bekannt? Ja/Nein). Ebenso wurde die Frage zum Influencer-Typ mittels einer 1-Item-Dichotomen-Skala erfasst. Das Geschlecht wurde über eine Auswahlbox mit den Optionen weiblich, männlich und divers erfragt, während das Alter in einem offenen Eingabefeld angegeben wurde.

Variable	Quelle	Dimension	Skala	α
Marken-Einstellung	Thomas & Fowler (2021) nach Till & Busler (2000)	Die Marke ist... gut/attraktiv/ positiv/ vorteilhaft	4-Item-7-Punkte-Semantisches Differential	0,98
Anzeigen-Einstellung	Franke et al. (2023) nach Silvera & Austad (2004)	Die Anzeige ist... angenehm/interessant/ sympathisch/gut	4-Item-7-Punkte-Semantisches Differential	0,92
Marken-Folgeabsicht	Adaptiert von Sands et al. (2022)	Wie wahrscheinlich ist es, dass Sie Marke X auf Instagram folgen würden?	1-Item-7-Punkte-Likert-Skala	-
Engagement-Rate	Eigene Arbeit	Würden Sie diesen Instagram-Beitrag liken/kommentieren/teilen?	3-Item-Dichotome-Skala	-
Kongruenz	Franke et al. (2023) nach Fleck et al. (2012)	Der Influencer und das beworbene Produkt... passen gut zusammen / harmonisieren gut / sind eine gute Kombination	3-Item-7-Punkte-Likert-Skala	0,85
Parasoziale Beziehung	Chung & Cho (2017)	Der Influencer gibt mir das Gefühl, als ob ich mit einem Freund zusammen wäre. ...	9-Item-7-Punkte-Likert-Skala	0,7

Tabelle 2: Operationalisierung der abhängigen Variablen und Mediatoren (Anmerkung: Die Fragen und Items wurden bereits ins Deutsche übersetzt.)

3.3 Experimentdesign

3.3.1 Forschungsdesign

Das Experiment folgt einem 2×2 Between-Subjects-Design mit den Faktoren Influencer-Typ (KI vs. Mensch) und Produktkategorie (Technik vs. Kosmetik). Der Influencer-Typ stellt die unabhängige Variable dar, während die Produktkategorie als Moderator wirkt. Daraus ergeben sich insgesamt die vier Untersuchungsgruppen: KI-Technik, KI-Kosmetik, Mensch-Technik und Mensch-Kosmetik. Die abhängigen Variablen umfassen die Marken-Einstellung, die Anzeigen-Einstellung, die Marken-Folgeabsicht und die Engagement-Rate. Kongruenz und parasoziale Beziehungen agieren als Mediatoren. Als Kontrollvariablen wurden das Produkt-Involvement (Franke et al., 2023), die technische Affinität, die Influencer-Bekanntheit (Franke et al., 2023), die Marken-Bekanntheit sowie Alter und Geschlecht berücksichtigt.

3.3.2 Stimulusmaterial

Influencer-Stimulus

In allen Gruppen wurde die britische, menschliche Influencerin Emily Bador (Bador, 2025) eingesetzt. Bador wurde aufgrund ihrer Ähnlichkeit zur bekannten KI-Influencerin Lil Miquela sowie ihrer Verwendung in früheren Studien (Song et al., 2024; Franke et al., 2023) ausgewählt. Um mögliche Störfaktoren wie Attraktivität, Stil und Ausdruck zu minimieren, wurden in allen Gruppen dasselbe Foto der Influencerin (Bador, 2022) und derselbe Postingtext verwendet. Um Vorerfahrungen oder Assoziationen mit ihrem realen Namen auszuschließen, erhielt sie einen fiktiven Namen (Franke et al., 2023). Dafür wurde Sarah Watson gewählt. Das Instagram-Posting, das als Stimulus diente und allen Probanden gezeigt wurde, ist in Abbildung 4 dargestellt. Der einzige Unterschied zwischen den Gruppen mit der KI-Influencerin und der menschlichen Influencerin bestand im Szenario-Text, der die Influencerin vorstellte. Die Szenario-Texte wurden basierend auf den Studien von Sands et al. (2022) und Thomas und Fowler (2021) entwickelt. Während die KI-Influencerin eine spezifische Einführung in ihren künstlichen Charakter erhielt, wurde dieser Abschnitt bei der menschlichen Influencerin ausgelassen. Beide Influencer-Typen erhielten jedoch denselben allgemeinen Szenario-Text.

Der Szenario-Text für die KI-Influencerin begann mit der folgenden Einführung: *Sarah Watson ist eine KI-Influencerin. Sie ist kein Mensch, sondern eine virtuelle Kreation, die ausschließlich auf Instagram existiert. Eine Agentur hat sie mithilfe von KI-Technologien erschaffen. Diese Technologien generieren Inhalte wie Fotos, Videos und Texte, die gezielt auf ihre Follower und Werbepartner abgestimmt sind. Die Agentur überprüft dabei die Inhalte, legt Themen und Ziele fest, gibt die Postings frei und schließt Werbekooperationen ab. Nun mehr zu ihrem Instagram-Auftritt:*

Anschließend folgte, für beide Influencer-Typen identisch, der allgemeine Szenario-Text: *Sarah Watson ist seit drei Jahren eine erfolgreiche Instagram-Influencerin mit über einer Million Followern. Sie promotet Produkte, teilt persönliche Einblicke und begeistert mit inspirierenden Beiträgen. Ihre Themen sind Mode, Kosmetik, Fitness und Tech-Trends.*

Produktkategorie-Stimulus

Der Moderator Produktkategorie wurde durch einen Instagram-Post für ein technisches und einen Instagram-Post für ein kosmetisches Produkt umgesetzt. Stellvertretend für die technische Kategorie wurde der Musik-Streaming-Dienst Spotify und für die kosmetische Kategorie der Lippenpflegestift Labello ausgewählt. Beide Marken sind im deutschsprachigen Raum bekannt, auf Instagram aktiv und kooperieren regelmäßig mit Influencern. Die Zuordnung zu den jeweiligen Produktkategorien ist eindeutig, da Spotify eine technische Softwarelösung darstellt, während Labello ein kosmetisches Pflegeprodukt ist.

Die Produkt-Postings wurden selbst erstellt und fotografiert. Dabei hielt eine weibliche Person das jeweilige Produkt in Richtung Himmel, während es fotografiert wurde. Es wurde auf einen neutralen Hintergrund sowie eine einheitliche Präsentation geachtet, um potenzielle Störfaktoren wie visuelle Unterschiede im Hintergrund oder in der Produktpräsentation zu minimieren. Der Instagram-Post für Spotify ist in Abbildung 5 und der für Labello in Abbildung 6 zu sehen und diente als Stimulusmaterial, das den jeweiligen Probanden in den entsprechenden Gruppen präsentiert wurde. Die grafische Bearbeitung und die Zusammenstellung der Postings, einschließlich der Integration von Fotos und Texten, erfolgten mithilfe der Software Figma (Figma, 2025) und eines dort bereitgestellten Smartphone-Mockups

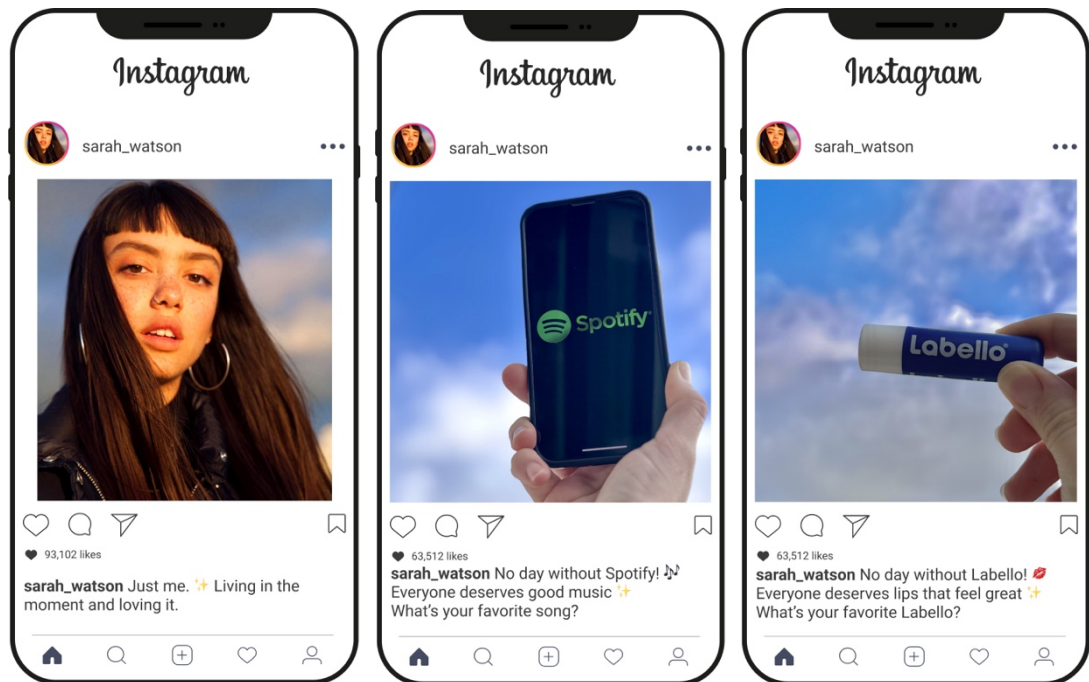


Abbildung 4: Influencer-Post

Abbildung 5: Spotify-Post

Abbildung 6: Labello-Post

3.3.3 Prozedur

Die Umfrage wurde mithilfe der Umfragesoftware SoSci Survey (SoSci Survey, 2025) durchgeführt. Der Ablauf der Umfrage wird in Abbildung 7 veranschaulicht. Zu Beginn erhielten die Probanden einen kurzen Einleitungstext, der die Studie sowie den Umfrageaufbau erläuterte, jedoch ohne das Thema KI zu erwähnen. Dies sollte mögliche Voreingenommenheiten vermeiden, insbesondere in der menschlichen Gruppe, da die Erwähnung von KI zu Irritationen führen könnte, wenn anschließend nur ein menschlicher Influencer gezeigt wird.

Anschließend folgten die demografischen Fragen (Alter, Geschlecht). Danach wurden die Probanden zufällig einer der beiden Gruppen (KI- oder menschliche Influencer-Gruppe) zugeordnet und erhielten den jeweiligen Szenario-Text angezeigt, der die Influencerin vorstellte. Im Anschluss wurde beiden Gruppen derselbe Instagram-Post der Influencerin (siehe Abbildung 4) präsentiert. In der KI-Gruppe erfolgte daraufhin der erste Manipulationscheck, um sicherzustellen, dass die Probanden den Szenario-Text gelesen und verstanden haben, dass es sich bei der Influencerin um eine KI-

Influencerin handelt. Danach wurden die Probanden zufällig einer der beiden Produktgruppen (Spotify-Gruppe oder Labello-Gruppe) zugewiesen. Sie bekamen daraufhin eine Ankündigung einer Werbekooperation zu sehen. Anschließend wurden der jeweilige Produkt-Post zu Spotify oder Labello (siehe Abbildungen 5 und 6) gezeigt. Unterhalb des Postings wurden die Fragen zur Engagement-Rate und zur Marken-Folgeabsicht direkt auf derselben Seite gestellt, da Instagram-Nutzer in der Realität den Beitrag ebenfalls vor sich haben, wenn sie liken, kommentieren oder teilen.

Es folgten Fragen zur Anzeigen- und Marken-Einstellung. Anschließend wurden die Probanden zur wahrgenommenen Kongruenz zwischen Influencer und Produkt befragt, bevor die Erhebung der parasozialen Beziehungen begann. Dazu erhielten sie einen Einleitungstext, in dem sie gebeten wurden, sich vorzustellen, dass sie der Influencerin bereits seit einem Jahr folgen. Dieser Zusatz diente der besseren Messung parasozialer Beziehungen, da diese erst durch wiederholte Kontakte entstehen.

Zum Schluss wurden die Kontrollfragen zu Produkt-Involvement, Marken-Bekanntheit, Influencer-Bekanntheit und technischer Affinität gestellt. Nach den Kontrollfragen wurde den Probanden ein Feld für Anmerkungen bereitgestellt, um ihnen die Möglichkeit zu geben, Feedback oder Kommentare zur Umfrage zu hinterlassen. In den KI-Gruppen wurde zusätzlich ein zweiter Manipulationscheck durchgeführt, bei dem gefragt wurde, ob die Probanden die Fragen unter der Annahme beantwortet haben, dass es sich um einen KI-Influencer handelt.

Abschließend wurden die Probanden auf einer Dankesseite über den Studienhintergrund informiert, einschließlich einer Erklärung zum Thema KI-Influencer. Vor dem Start der Hauptbefragung wurde ein Pretest mit insgesamt zwölf Personen (vier pro Gruppe) durchgeführt, um Verständnisprobleme und Schwächen im Fragebogen aufzudecken. Der Pretest verlief erfolgreich, sodass die Umfrage wie geplant starten konnte.

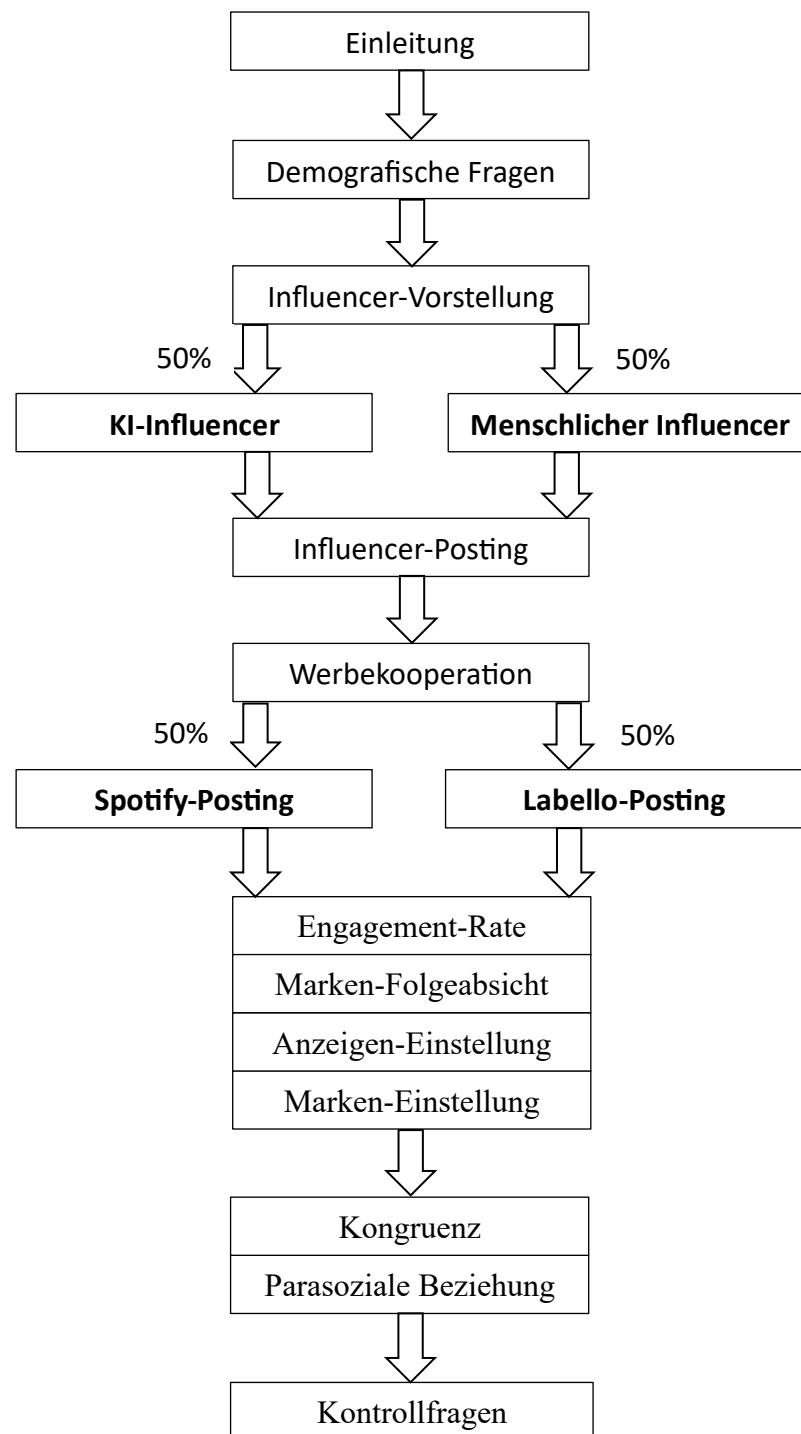


Abbildung 7: Ablaufplan des Experiments

4. Ergebnisse

In diesem Kapitel werden die erhobenen Daten analysiert und die Hypothesen überprüft. Zunächst erfolgt im Kapitel *Variablenkonstruktion und Testvoraussetzungen* eine Beschreibung der Konstruktion der Messvariablen, die Bestimmung der Skalenniveaus und die Analyse der Reliabilitätsmaße. Darüber hinaus werden die Voraussetzungen für die parametrischen Tests geprüft, darunter die Normalverteilung, Schiefe und Homogenität der Varianzen. Im Kapitel *Deskriptive Statistik* wird ein Überblick über die Daten und deren Struktur gewonnen, indem zentrale Kennwerte wie Mittelwert, Median und Standardabweichung ermittelt werden. Die Daten werden zusätzlich durch Boxplots visualisiert und Zusammenhänge zwischen den Variablen werden mithilfe einer Korrelationstabelle dargestellt. Im Kapitel *Schließende Statistik* werden die Hypothesen H1 bis H9 mithilfe geeigneter statistischer Verfahren getestet und die Ergebnisse dargestellt.

4.1 Variablenkonstruktion und Testvoraussetzungen

Für die Analyse und Auswertung der Daten wurde das Statistikprogramm R (R Core, 2024) verwendet. Die Messvariablen wurden entsprechend ihres Skalenniveaus klassifiziert. Geschlecht, Produktbekanntheit und Influencer-Bekanntheit sind nominale Variablen, während Alter als metrische Variable erfasst wurde. Die Variablen Marken-Einstellung, Anzeigen-Einstellung, Kongruenz und parasoziale Beziehungen wurden als metrisch interpretiert. Die Werte der einzelnen Items, die auf einer 7-Punkte-Likert-Skala oder einem 7-Punkte-Semantischen Differential basierten, wurden aufsummiert und durch die Anzahl der Items dividiert. Obwohl Likert-Skalen theoretisch nur ordinale Daten liefern, wird davon ausgegangen, dass die Antwortoptionen gleichmäßige Intervallschritte suggerieren und somit eine Interpretation als metrische Variablen zulassen. Dieses Vorgehen entspricht den methodischen Standards aktueller wissenschaftlicher Arbeiten (Thomas & Fowler, 2021; Sands et al., 2022; Franke et al., 2023). Die Variablen Marken-Folgeabsicht, Produkt-Involvement und technische Affinität wurden jeweils durch ein einzelnes 7-Punkte-Item erhoben und als metrisch eingestuft. Dies entspricht dem Vorgehen von Franke et al. (2023), die das Produkt-Involvement, und von Sands et al. (2022), die die Influencer-Folgeabsicht auf die gleiche Weise erfasst und bewertet haben.

Die Variable Engagement-Rate wird grundsätzlich aggregiert berechnet, indem die Summe aller Interaktionen (Likes, Kommentare und Teilungen) durch die Anzahl der Probanden geteilt und das Ergebnis mit 100 multipliziert wird, um einen Prozentwert zu erhalten. Auf Instagram erfolgt diese Berechnung typischerweise im Verhältnis zur Anzahl der Follower (Piwik, 2025). Diese aggregierte Berechnung liefert für jede Untersuchungsgruppe lediglich einen einzigen Prozentwert bzw. Datenpunkt, wodurch die Varianz und Häufigkeitsverteilung innerhalb der Gruppen verloren geht und Hypothesentests nicht durchführbar wären. Um dies zu vermeiden und alle Datenpunkte für die statistischen Analysen beizubehalten, wurde ein individueller Engagement-Score berechnet. Jedem Probanden wurde dabei ein Score zugewiesen, der der Summe seiner Interaktionen entspricht. So ergeben beispielsweise ein Like und ein Kommentar zusammen zwei Punkte. Dieser Score weist Ausprägungen von 0 (keine Interaktion) bis 3 (Like, Kommentar, Teilen) auf, besitzt Varianz und ist ordinal skaliert. Dies ermöglicht den Einsatz von Verfahren wie Chi²-Tests zur Untersuchung von Gruppenunterschieden und ordinaler Regressionsmodelle zur Analyse von Zusammenhängen. Die aggregierte Engagement-Rate wird weiterhin für deskriptive Gegenüberstellungen verwendet, während die statistischen Tests auf dem Engagement-Score basieren.

Um die Qualität der verwendeten Messinstrumente sicherzustellen, wurde die Reliabilität der gebildeten Skalen durch Cronbachs Alpha geprüft. Alle Skalen wiesen Werte über dem empfohlenen Schwellenwert von 0,70 auf, was auf eine hohe Reliabilität der Messinstrumente deutet (Anzeigen-Einstellung: $\alpha = 0,86$; Marken-Einstellung: $\alpha = 0,95$; Kongruenz: $\alpha = 0,94$; parasoziale Beziehungen: $\alpha = 0,96$).

Für die geplanten Hypothesentests müssen bestimmte methodische Voraussetzungen erfüllt sein, da in der Analyse parametrische Verfahren eingesetzt werden. Insbesondere ist es erforderlich, die Normalverteilung der Daten, die Homogenität der Varianzen sowie die Schiefe zu überprüfen. Gemäß dem Zentralen Grenzwertsatz (CLT) sind Mittelwerte von Stichproben bei ausreichend großen Stichprobengrößen annähernd normalverteilt, unabhängig von der Verteilung der Ausgangsdaten. Da in jeder Gruppe mindestens 50 Probanden vorhanden sind, ist diese Bedingung für parametrische Tests erfüllt. Dennoch wurden die metrischen Variablen zusätzlich mithilfe des Lilliefors-Tests und visueller Inspektionen durch QQ-Plots auf

Normalverteilung geprüft. Während die Daten der Anzeigen-Einstellung in allen Gruppen die Normalverteilungsannahme erfüllten, zeigten andere Variablen kleine bis mittlere Abweichungen. Die Levene-Tests zur Überprüfung der Varianzhomogenität verliefen bei allen Variablen erfolgreich. Eine Analyse der Schiefe ergab jedoch, dass die Variable Marken-Folgeabsicht eine starke Rechtsschiefe aufweist (Schiefe = 1,7), was darauf zurückzuführen ist, dass viele Probanden niedrige Werte (1 oder 2 auf einer 7-Punkte-Skala) angaben und zugleich starke Ausreißer nach oben hin enthalten waren. Zur Korrektur wurde diese Variable logarithmisch transformiert, wodurch die Schiefe reduziert bzw. abgedämpft wurde und die Voraussetzung für parametrische Tests sichergestellt wurde. Die Interpretation und Visualisierung (Boxplot) der Folgeabsicht erfolgt weiterhin auf Basis der Originalwerte, während die statistischen Tests mit der logarithmierten Version durchgeführt werden.

4.2 Deskriptive Statistik

Die Mittelwerte und Standardabweichungen der abhängigen Variablen (Marken-Einstellung, Anzeigen-Einstellung, Marken-Folgeabsicht, Engagement-Rate) und der Mediatoren (Kongruenz und parasoziale Beziehungen) sind in Tabelle 3 dargestellt.

Die Ergebnisse zeigen, dass menschliche Influencer innerhalb einer Produktkategorie durchgängig bessere Ergebnisse erzielen als KI-Influencer. Die größten Unterschiede zeigen sich bei der Engagement-Rate, der Kongruenz und den parasozialen Beziehungen, während die Marken-Einstellung und die Anzeigen-Einstellung moderate Differenzen aufweisen. Zudem zeigt sich, dass alle untersuchten Variablen bei beiden Influencer-Typen in der Technik-Kategorie (Spotify) höher ausfallen als in der Kosmetik-Kategorie (Labello). Eine Ausnahme bilden die Anzeigen-Einstellung bei menschlichen Influencern und die parasozialen Beziehungen bei KI-Influencern, bei denen die Werte in beiden Kategorien nahezu gleich sind. Insgesamt wird deutlich, dass KI-Influencer nur dann besser als menschliche Influencer abschneiden, wenn bei den markenbezogenen Variablen (Marken-Einstellung und Marken-Folgeabsicht) der Vergleich zwischen KI-Influencern in der Technik-Kategorie und menschlichen Influencern in der Kosmetik-Kategorie gezogen wird. In allen anderen Vergleichen erzielen menschliche Influencer höhere Werte. Zur weiteren Analyse werden die Werte und Verteilungen der einzelnen Variablen mithilfe von Boxplots visualisiert

und gruppenweise verglichen. Ergänzend wird eine Korrelationstabelle erstellt, um mögliche Zusammenhänge zwischen den Variablen zu erkennen.

Gruppe/Variable	Marken-Einstellung (M ± SD)	Anzeigen-Einstellung (M ± SD)	Marken-Folgeabsicht (M ± SD)	Engagement-Rate %	Kongruenz (M ± SD)	Parasoziale Beziehung (M ± SD)
Mensch-Technik (n = 61)	5,66 ± 1,13	3,58 ± 1,21	2,39 ± 1,53	22,95%	4,56 ± 1,36	3,50 ± 1,39
KI-Technik (n = 56)	5,38 ± 1,33	3,25 ± 1,08	2,25 ± 1,24	12,50%	3,88 ± 1,62	2,59 ± 1,26
Mensch-Kosmetik (n = 61)	4,81 ± 0,96	3,65 ± 1,30	1,93 ± 1,08	9,83%	4,53 ± 1,45	3,39 ± 1,52
KI-Kosmetik (n = 58)	4,79 ± 1,36	3,18 ± 1,26	1,83 ± 1,24	8,62%	3,66 ± 1,79	2,61 ± 1,32
Mensch-Gesamt (n=122)	5,23 ± 1,13	3,61 ± 1,25	2,16 ± 1,34	16,39%	4,54 ± 1,40	3,44 ± 1,45
KI-Gesamt (n=114)	5,08 ± 1,37	3,21 ± 1,17	2,04 ± 1,25	10,53%	3,76 ± 1,70	2,60 ± 1,28
Technik-Gesamt (n=117)	5,52 ± 1,23	3,42 ± 1,16	2,32 ± 1,39	17,95%	4,23 ± 1,52	3,06 ± 1,40
Kosmetik-Gesamt (n=119)	4,80 ± 1,17	3,42 ± 1,30	1,88 ± 1,16	9,24%	4,10 ± 1,68	3,01 ± 1,48

Tabelle 3: Zusammenfassung der deskriptiven Werte (M: Mittelwert; SD: Standardabw.; 7-Punkte-Likert-Skala); KI-Gesamt = KI-Kosmetik + KI-Technik; Mensch-Gesamt = Mensch-Kosmetik + Mensch-Technik; Technik-Gesamt: KI-Technik + Mensch-Technik; Kosmetik-Gesamt: KI-Kosmetik + Mensch-Kosmetik

Marken-Einstellung

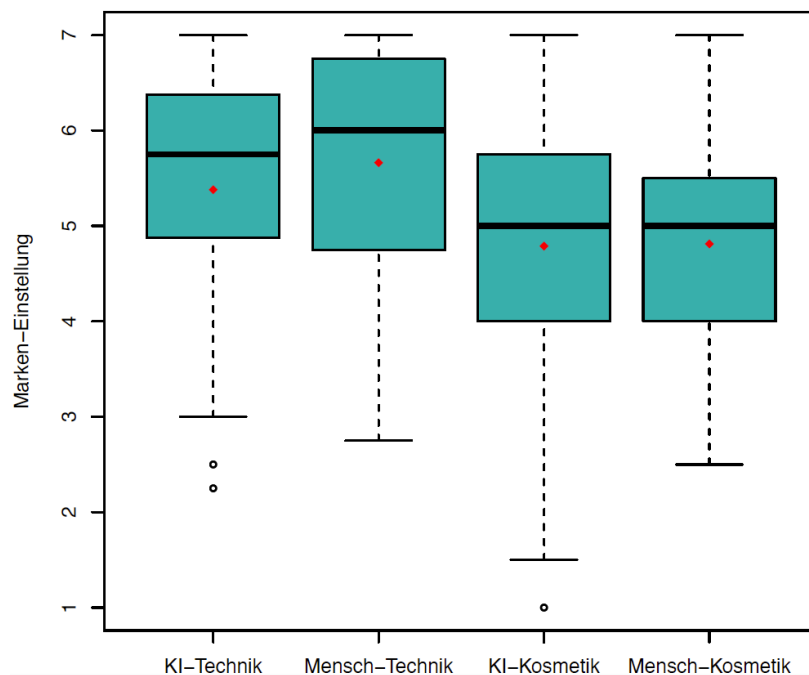


Abbildung 8: Boxplot Marken-Einstellung

Die Marken-Einstellung ist höher, wenn mit menschlichen Influencern geworben wird als mit KI-Influencern, wobei der Unterschied in der Technik-Kategorie ausgeprägter ist. Insgesamt wird die Marken-Einstellung in der Technik-Kategorie deutlich positiver bewertet. In dieser Kategorie beträgt der Mittelwert für menschliche Influencer $M = 5,66$ ($SD = 1,13$) und für KI-Influencer $M = 5,38$ ($SD = 1,33$). In der Kosmetik-Kategorie liegen die Mittelwerte bei $M = 4,81$ ($SD = 0,96$) für menschliche und $M = 4,79$ ($SD = 1,36$) für KI-Influencer. Es zeigt sich somit eine insgesamt positive Einstellung gegenüber den Marken. Die Boxplots (siehe Abbildung 8) verdeutlichen diese Ergebnisse. Die Mittelwerte sind als rote Punkte dargestellt und liegen nahe den jeweiligen Medianen, was auf eine überwiegend symmetrische Verteilung hindeutet. Bei den KI-Influencern treten vereinzelt kleinere Ausreißer nach unten auf.

Anzeigen-Einstellung

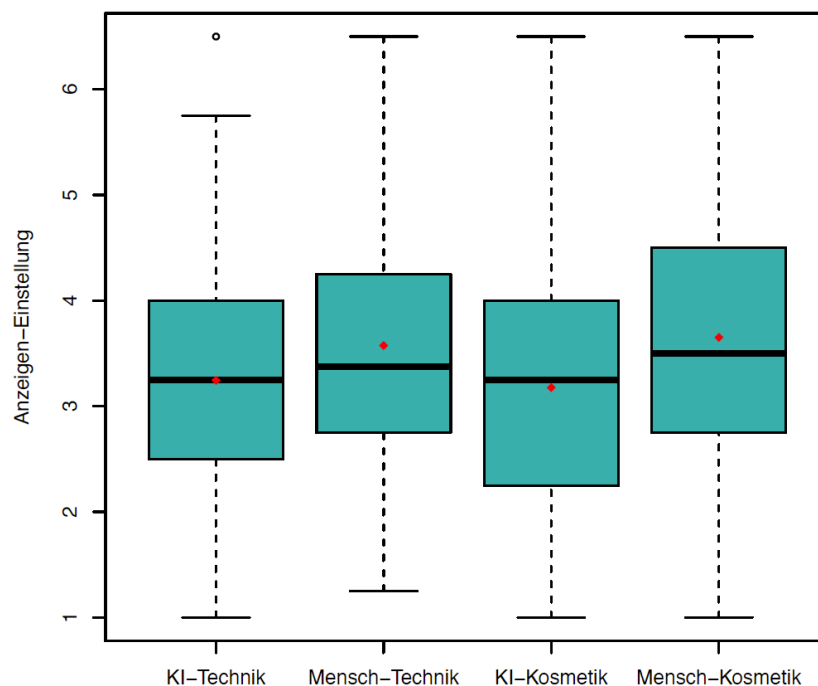


Abbildung 9: Boxplot-Anzeigen-Einstellung

Menschliche Influencer erzielen durchgängig eine höhere Anzeigen-Einstellung als KI-Influencer. Während menschliche Influencer in der Kosmetik-Kategorie höhere Werte erreichen als in der Technik-Kategorie, bleiben die Werte bei KI-Influencern in beiden Kategorien auf einem ähnlichen Niveau, mit einer geringfügigen Erhöhung in der Technik-Kategorie. In der Technik-Kategorie beträgt der Mittelwert für

menschliche Influencer $M = 3,58$ ($SD = 1,21$) und für KI-Influencer $M = 3,25$ ($SD = 1,08$). In der Kosmetik-Kategorie liegt der Mittelwert bei menschlichen Influencern bei $M = 3,65$ ($SD = 1,30$) und bei KI-Influencern bei $M = 3,18$ ($SD = 1,26$). Insgesamt zeigt sich eine neutrale (menschliche Influencer) bis leicht negative (KI-Influencer) Anzeigen-Einstellung. Die Boxplots (siehe Abbildung 9) zeigen, dass Mittelwerte und Mediane nahe beieinander liegen, was auf eine symmetrische Verteilung der Daten hindeutet.

Marken-Folgeabsicht

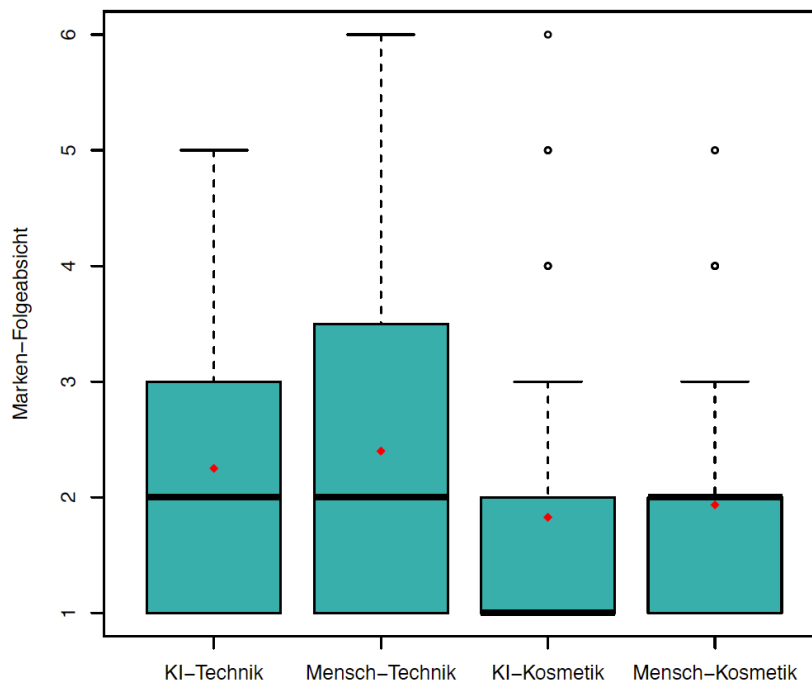


Abbildung 10: Boxplot Marken-Folgeabsicht

Menschliche Influencer weisen innerhalb einer Produktkategorie eine geringfügig höhere Marken-Folgeabsicht auf als KI-Influencer, wobei beide Influencer-Typen in der Technik-Kategorie deutlich höhere Werte erreichen. In der Technik-Kategorie beträgt der Mittelwert für menschliche Influencer $M = 2,39$ ($SD = 1,53$) und für KI-Influencer $M = 2,25$ ($SD = 1,24$). In der Kosmetik-Kategorie liegt der Mittelwert bei menschlichen Influencern bei $M = 1,93$ ($SD = 1,08$), während KI-Influencer $M = 1,83$ ($SD = 1,24$) erreichen. Die insgesamt niedrigen Mittelwerte deuten jedoch auf eine geringe Bereitschaft zur Marken-Folgeabsicht hin. Die Boxplots (siehe Abbildung 10) bestätigen diese Tendenz. Das 25 %-Quartil liegt bei 1, beim KI-Influencern in der Kosmetik-Kategorie befindet sich der Median sogar bei 1. Die starke Rechtsschiefe und das Vorhandensein von Ausreißern bei den Kosmetik-Kategorien stützen die

Entscheidung, die Skala für die statistischen Tests logarithmisch zu transformieren, um eine symmetrischere Verteilung zu erreichen.

Engagement-Rate

Die Engagement-Raten sind durchgängig hoch, wobei die Interaktionen mit menschlichen Influencern häufiger ausfielen. In der Technik-Kategorie beträgt die Engagement-Rate für menschliche Influencer 22,95 % (12 Likes, 1 Kommentar, 1 geteilte Interaktion bei 61 Probanden) und für KI-Influencer 12,5 % (6 Likes, 1 Kommentar bei 56 Probanden). In der Kosmetik-Kategorie liegt die Engagement-Rate für menschliche Influencer bei 9,83 % (6 Likes bei 61 Probanden) und für KI-Influencer bei 8,62 % (5 Likes bei 58 Probanden).

Visualisierung: Da die Interaktionen nahezu ausschließlich in Form von Likes erfolgten und bei einer Interaktion in der Regel ein Engagement-Score von 1 erzielt wurde und weil die Engagement-Rate selbst keine Varianz aufweist, wurde auf eine grafische Darstellung der Engagement-Raten bzw. der Engagement-Scores verzichtet.

Kongruenz

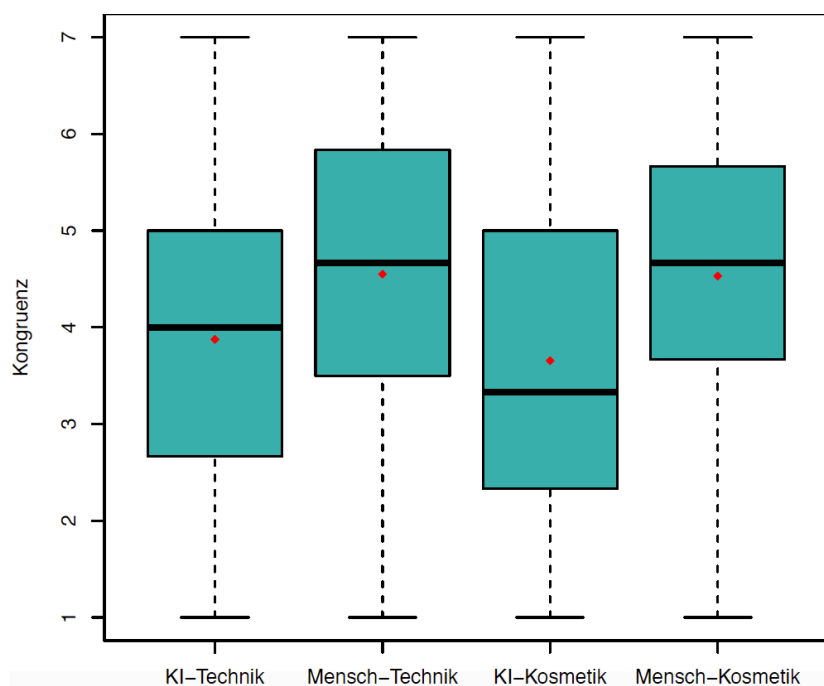


Abbildung 11: Boxplot Kongruenz

Die wahrgenommene Kongruenz zwischen dem Influencer und dem beworbenen Produkt ist bei menschlichen Influencern durchgängig deutlich höher als bei KI-

Influencern. Während die Kongruenzwerte bei menschlichen Influencern in beiden Produktkategorien nahezu konstant bleiben, fällt bei KI-Influencern in der Technik-Kategorie eine leichte Erhöhung auf. In der Technik-Kategorie beträgt der Mittelwert für menschliche Influencer $M = 4,56$ ($SD = 1,36$) und für KI-Influencer $M = 3,88$ ($SD = 1,62$). In der Kosmetik-Kategorie liegen die Mittelwerte für menschliche Influencer bei $M = 4,53$ ($SD = 1,45$) und für KI-Influencer bei $M = 3,66$ ($SD = 1,79$). Insgesamt bewegen sich die Kongruenzwerte im neutralen (menschliche Influencer) bis negativen Bereich (KI-Influencer). Die Medianwerte entsprechen weitgehend den Mittelwerten (siehe Abbildung 11), was auf eine symmetrische Verteilung hinweist.

Parasoziale Beziehung

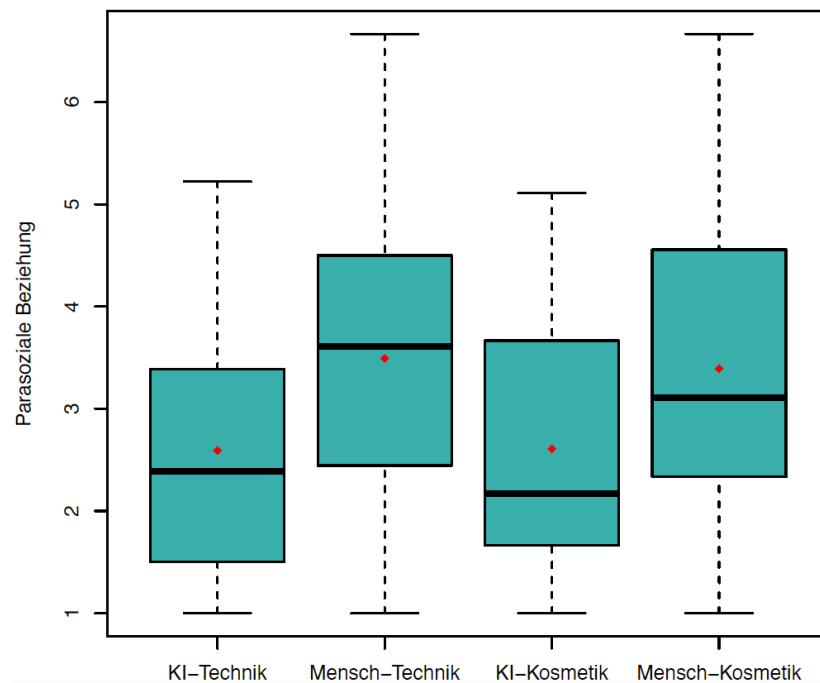


Abbildung 12: Boxplot Parasoziale Beziehung

Die parasoziale Beziehung zu menschlichen Influencern ist durchgängig deutlich stärker ausgeprägt als zu KI-Influencern. Die Intensität bleibt über die Produktkategorien hinweg weitgehend konstant. In der Technik-Kategorie beträgt der Mittelwert für menschliche Influencer $M = 3,50$ ($SD = 1,39$) und für KI-Influencer $M = 2,59$ ($SD = 1,26$). In der Kosmetik-Kategorie liegen die Werte bei $M = 3,39$ ($SD = 1,52$) für menschliche Influencer und $M = 2,61$ ($SD = 1,32$) für KI-Influencer. Insgesamt sind die parasozialen Beziehungen bei den KI-Influencern schwach und bei den menschlichen Influencern stärker ausgeprägt. Die Medianwerte liegen nahe den Mittelwerten (siehe Abbildung 12), was auf eine symmetrische Verteilung hinweist.

Korrelationen

Die Korrelationen zeigen mehrere Zusammenhänge zwischen den untersuchten Variablen. Die Mediatoren Kongruenz und parasoziale Beziehung korrelieren sowohl untereinander ($r = 0,42$) als auch mit der Anzeigen-Einstellung (Kongruenz: $r = 0,53$; parasoziale Beziehung: $r = 0,48$). Dies deutet darauf hin, dass beide Variablen mit der Anzeigen-Einstellung in Zusammenhang stehen. Zudem besteht eine positive Korrelation zwischen der Anzeigen-Einstellung und der Marken-Folgeabsicht ($r = 0,46$) und der Marken-Einstellung ($r = 0,41$), was darauf hindeutet, dass eine positive Anzeigen-Einstellung auch die Marken-Einstellung und Marken-Folgeabsicht beeinflussen könnte. Die Kontrollvariable Produkt-Involvement zeigt eine hohe Korrelation mit der Marken-Einstellung ($r = 0,60$), eine moderate Korrelation mit der Marken-Folgeabsicht ($r = 0,39$) und geringere Korrelationen mit der Anzeigen-Einstellung ($r = 0,26$) und den parasozialen Beziehungen ($r = 0,22$), was auf eine mögliche Kovariate hindeutet.

Variable	Marken-Einstellung	Anzeigen-Einstellung	Marken-Folgeabsicht	Kongruenz	Parasoziale Beziehung	Produkt-Involvement
Marken-Einst.	1,00	0,41	0,30	0,08	0,19	0,60
Anzeigen-Einst.	0,41	1,00	0,46	0,53	0,48	0,26
Marken-Folgeabs.	0,30	0,46	1,00	0,24	0,34	0,39
Kongruenz	0,08	0,53	0,24	1,00	0,42	0,05
Parasoziale Bez.	0,19	0,48	0,34	0,42	1,00	0,22
Produkt-Involv.	0,60	0,26	0,39	0,05	0,22	1,00

Tabelle 4: Korrelationstabelle; Anmerkung: Die Engagement-Rate wurde in der Korrelationstabelle nicht berücksichtigt, da sie im aggregierten Zustand jeweils nur einen einzelnen Wert aufweist und als Engagement-Score nicht metrisch skaliert ist.

4.3 Schließende Statistik

Bevor die Hypothesen H1 bis H9 statistisch geprüft werden, wurde untersucht, ob sich die vier Gruppen (KI-Technik, Mensch-Technik, KI-Kosmetik, Mensch-Kosmetik) hinsichtlich der Kontrollvariablen Alter, Geschlecht, Produkt-Involvement und technischer Affinität signifikant unterscheiden. Dies dient dazu, sicherzustellen, dass die Testergebnisse möglichst frei von Störeinflüssen sind bzw. dass diese kontrolliert werden können.

Die ANOVA für das Alter ergab keinen signifikanten Unterschied zwischen den Gruppen ($F(3;232) = 1,19$; $p = 0,315$). Der Chi-Quadrat-Test für das Geschlecht zeigte ebenfalls keinen signifikanten Unterschied ($\chi^2 = 3,38$; $df = 6$; $p = 0,76$). Für die technische Affinität ergab die ANOVA ebenfalls keine signifikanten Gruppenunterschiede ($F(3;232) = 1,71$; $p = 0,165$). Ein signifikanter Unterschied zeigte sich hingegen beim Produkt-Involvement. Dabei wurden in der Technik-Kategorie signifikant höhere Werte erzielt ($F(1;232) = 86,78$; $p < 0,001$). Der Influencer-Typ hatte jedoch keinen Einfluss auf das Produkt-Involvement ($F(1;232) = 0,004$; $p = 0,950$), ebenso wenig die Interaktion ($F(1,232) = 0,93$; $p = 0,336$). Da sich das Produkt-Involvement signifikant zwischen den Produktkategorien unterscheidet, Korrelationen mit den abhängigen Variablen und den Mediatoren (mit Ausnahme der Kongruenz) bestehen und es von Franke und Groeppel-Klein (2024) ebenfalls in den Tests kontrolliert wurde, wird es in den nachfolgenden Analysen als Kovariate einbezogen, um potenzielle Störeinflüsse auszuschließen.

Zur Sicherstellung einer ausreichenden statistischen Power wurde vorab eine Power-Analyse durchgeführt. Mit einer geschätzten Effektgröße von $f = 0,2$ (kleiner Effekt), einem Signifikanzniveau von 0,05 und einer Power von 0,80 wäre eine Mindeststichprobengröße von 198 Probanden erforderlich gewesen. Mit 236 Probanden liegt die tatsächliche Stichprobengröße darüber, sodass eine ausreichend große statistische Power gewährleistet werden kann. Die Annahme eines kleinen Effekts basiert auf früheren Studien zu KI-Influencern (Franke et al., 2023; Sands et al., 2022).

Nun folgt die statistische Überprüfung der Hypothesen H1 bis H9 (alle relevanten Mittelwerte inkl. Standardabweichungen sind in Tabelle 3 zu finden):

H1: Es wurde geprüft, ob KI-Influencer im Vergleich zu menschlichen Influencern eine höhere Marken-Folgeabsicht (H1a) und Engagement-Rate (H1b) erzielen. Gleichzeitig wurde erwartet, dass bei der Marken-Einstellung (H1c) vergleichbare Werte erzielt werden und die Anzeigen-Einstellung (H1d) bei KI-Influencern geringer ausfällt. Zur Kontrolle des Einflusses wurde das Produkt-Involvement als Kovariate in die ANCOVA-Modelle (H1a, H1b, H1d) einbezogen, wobei es in diesen Analysen

jeweils einen signifikanten Einfluss zeigte (alle $p < 0,001$). Die logarithmierte Marken-Folgeabsicht (H1a) zeigte keinen signifikanten Unterschied zwischen den Influencer-Typen ($F(1; 233) = 0,671$; $p = 0,414$). H1a wurde somit nicht bestätigt. Auch die Analyse der Engagement-Rate (H1b) mittels Chi-Quadrat-Test ergab keine signifikanten Unterschiede ($\chi^2(3) = 3,55$; $p = 0,314$), sodass H1b ebenfalls nicht bestätigt wurde. Die Marken-Einstellung (H1c) zeigte vergleichbare Werte zwischen den Influencer-Typen ($F(1;233) = 1,396$; $p = 0,24$), was die H1c stützt. Die Analyse der Anzeigen-Einstellung (H1d) ergab einen signifikanten Unterschied ($F(1;233) = 6,973$; $p = 0,009$), wobei menschliche Influencer ($M = 3,61$; $SD = 1,25$) höhere Werte erzielten als KI-Influencer ($M = 3,21$; $SD = 1,17$), weshalb die Hypothese H1d bestätigt werden konnte. Zusammenfassend konnten H1a und H1b nicht bestätigt werden, während H1c und H1d bestätigt wurden.

H2: Es wurde getestet, ob die Produktkategorie (Technik vs. Kosmetik) die Wirkung von Influencern moderiert, also ob KI-Influencer bei technischen Produkten bessere Ergebnisse erzielen als bei kosmetischen Produkten und ob menschliche Influencer bei kosmetischen Produkten bessere Ergebnisse erzielen als bei technischen Produkten. Dies wurde hinsichtlich der Marken-Einstellung (H2a), der Anzeigen-Einstellung (H2b), der Marken-Folgeabsicht (H2c) und der Engagement-Rate (H2d) untersucht. Das Produkt-Involvement wurde als Kovariate in die ANCOVA-Modelle (H2a, H2b, H2c) eingefügt und zeigte jeweils einen signifikanten Einfluss (alle $p < 0,001$). Bei der Marken-Einstellung (H2a) ergab der Test einen signifikanten Effekt der Produktkategorie, wobei in der Technik-Kategorie ($M = 5,52$; $SD = 1,23$) höhere Bewertungen erzielt wurden als in der Kosmetik-Kategorie ($M = 4,80$; $SD = 1,17$) ($F(1;231) = 30,23$; $p < 0,001$) und einen nicht signifikanten Effekt des Influencer-Typs ($F(1;231) = 1,39$; $p = 0,24$). Die Interaktion war nicht signifikant ($F(1;231) = 0,12$; $p = 0,735$), somit konnte H2a nicht bestätigt werden. Bei der Anzeigen-Einstellung (H2b) zeigte der Test wiederholt einen signifikanten Effekt des Influencer-Typs ($F(1;231) = 7,13$; $p = 0,008$) und einen nicht signifikanten Effekt der Produktkategorie ($F(1;231) = 0,001$; $p = 0,97$). Die Interaktion war nicht signifikant ($F(1;231) = 0,62$; $p = 0,43$), somit konnte H2b nicht bestätigt werden. Bei der logarithmierten Marken-Folgeabsicht (H2c) ergab der Test einen signifikanten Effekt der Produktkategorie, wobei in der Technik-Kategorie ($M = 2,32$; $SD = 1,39$) höhere Bewertungen erzielt wurden als in der Kosmetik-Kategorie ($M =$

1,88; SD = 1,16) ($F(1;231) = 7,79$; $p = 0,006$) und einen nicht signifikanten Effekt des Influencer-Typs ($F(1;231) = 0,67$; $p = 0,42$). Die Interaktion war nicht signifikant ($F(1;231) = 0,86$; $p = 0,35$), somit konnte H2c nicht bestätigt werden. Der Chi-Quadrat-Test der Engagement-Rate (H2d) zeigte ebenfalls keine moderierende Wirkung ($\chi^2(9) = 9,77$; $p = 0,37$), sodass H2d nicht bestätigt werden konnte. Insgesamt konnten H2a, H2b, H2c und H2d nicht bestätigt werden.

H3: Es wurde geprüft, ob KI-Influencer und menschliche Influencer eine ähnliche Kongruenz erzielen, wenn die Produktkategorie nicht einbezogen wird. Zur Kontrolle wurde das Produkt-Involvement als Kovariate in ein ANCOVA-Modell aufgenommen, dabei ergab sich kein signifikanter Effekt ($p = 0,41$). Die Ergebnisse ergaben einen signifikanten Unterschied zwischen den Influencer-Typen ($F(1;233) = 14,82$; $p < 0,001$), der menschliche Influencer ($M = 4,54$; $SD = 1,40$) erzielte dabei eine höhere Kongruenz als der KI-Influencer ($M = 3,76$; $SD = 1,70$). Somit konnte die H3, wonach beide Influencer-Typen eine vergleichbare Kongruenz erzielen, nicht bestätigt werden.

H4: Es wurde geprüft, ob die Kongruenz zwischen Influencern und Produkt durch die Produktkategorie moderiert wird, also ob KI-Influencer eine höhere Kongruenz mit technischen Produkten und menschliche Influencer eine höhere Kongruenz mit kosmetischen Produkten erzielen. Zur Kontrolle wurde das Produkt-Involvement als Kovariate in das ANCOVA-Modell einbezogen, es zeigte sich kein signifikanter Einfluss ($p = 0,55$). Die Ergebnisse zeigten wiederholt einen signifikanten Effekt des Influencer-Typs ($F(1;231) = 14,716$; $p < 0,001$) und keinen signifikanten Effekt der Produktkategorie ($F(1;231) = 0,35$; $p = 0,55$). Der Interaktionseffekt war nicht signifikant ($F(1;231) = 0,26$; $p = 0,60$), somit konnte H4 nicht bestätigt werden.

H5: Es wurde untersucht, ob eine höhere Kongruenz zwischen Influencern und Produkt zu einer positiveren Anzeigen-Einstellung führt. Zur Überprüfung wurde ein lineares Regressionsmodell geschätzt, in dem Kongruenz als Prädiktor und Produkt-Involvement als Kovariate einbezogen wurden ($p < 0,001$). Die Analyse ergab einen signifikanten positiven Effekt der Kongruenz ($\beta = 0,401$; $t(233) = 9,79$; $p < 0,001$). Dies bedeutet, dass eine Erhöhung der Kongruenz um eine Einheit mit einem Anstieg der Anzeigen-Einstellung um 0,401 verbunden ist. Das Produkt-Involvement zeige ebenfalls einen signifikanten Effekt ($\beta = 0,149$; $t(233) = 4,35$; $p < 0,001$). Die

Residuenplots zeigten kein systematisches Muster, die Normalverteilung der Residuen wurde durch den Shapiro-Wilk-Test bestätigt und die Variance Inflation Factor (VIF)-Analyse ergab keine Hinweise auf Multikollinearität. Somit wurde H5 bestätigt.

Mediatoranalyse (Kongruenz): Zusätzlich wurde eine Mediationsanalyse durchgeführt, um zu überprüfen, ob der Effekt des Influencer-Typs auf die Anzeigen-Einstellung durch die Kongruenz vermittelt wird. Die Mediationsanalyse wurde mithilfe des PROCESS-Makros (Modell 4) für R von Hayes (2022) durchgeführt. Der Influencer-Typ (KI vs. Mensch) wurde als unabhängige Variable (X), die Kongruenz als Mediator (M) und die Anzeigen-Einstellung als abhängige Variable (Y) ins Modell eingegeben. Die Ergebnisse ergaben, dass der Influencer-Typ einen signifikanten Einfluss auf die Kongruenz hatte (Pfad a: $b = 0,78$; $t(234) = 3,85$; $p < 0,001$). Die Kongruenz wiederum hatte einen signifikanten positiven Effekt auf die Anzeigen-Einstellung (Pfad b: $b = 0,40$; $t(233) = 9,21$; $p < 0,001$). Der indirekte Effekt ($a \times b$) war signifikant ($b = 0,315$; $[0,14, 0,50]$), während der direkte Effekt des Influencer-Typs auf die Anzeigen-Einstellung nicht signifikant war (Pfad c': $b = 0,089$; $t(233) = 0,64$; $p = 0,525$). Der totale Effekt war hingegen signifikant (Pfad c: $b = 0,402$; $p = 0,008$). 76,42 % des Effekts wurden über die Kongruenz vermittelt. Dies bestätigt, dass der Einfluss des Influencer-Typs auf die Anzeigen-Einstellung durch die wahrgenommene Kongruenz vollständig vermittelt wird.

H6: Es wurde untersucht, ob die Intensität der parasozialen Beziehung bei menschlichen Influencern höher ist als bei KI-Influencern. Zur Kontrolle des Einflusses wurde das Produkt-Involvement ($p < 0,001$) als Kovariate in das ANCOVA-Modell einbezogen. Die Analyse ergab einen signifikanten Effekt des Influencer-Typs ($F(1;233) = 23,45$; $p < 0,001$). Menschliche Influencer ($M = 3,44$; $SD = 1,45$) weisen eine signifikant stärkere parasoziale Beziehung als KI-Influencer ($M = 2,60$; $SD = 1,28$) auf, sodass H6 bestätigt wurde. Zusätzlich wurde explorativ getestet, ob die Produktkategorie die parasoziale Beziehung moderiert. Die Ergebnisse zeigen, dass weder die Produktkategorie noch die Interaktion signifikant waren.

H7: Es wurde getestet, ob eine höhere parasoziale Beziehung zu einer höheren Anzeigen-Einstellung führt. Dazu wurde ein lineares Regressionsmodell geschätzt, in dem die parasoziale Beziehung als Prädiktor und das Produkt-Involvement als

Kovariate einbezogen wurden ($p = 0,006$). Die Analyse ergab einen signifikant positiven Effekt der parasozialen Beziehung ($\beta = 0,383$; $t(233) = 7,713$; $p < 0,001$) sowie einen signifikanten Effekt des Produkt-Involvements ($\beta = 0,103$; $t(233) = 2,765$; $p = 0,006$). Die Annahmen der Normalverteilung und Homoskedastizität der Residuen sowie das Fehlen von Multikollinearität wurden erfolgreich geprüft. Somit wurde H7 bestätigt.

Mediatoranalyse (parasoziale Beziehung): Zur weiteren Überprüfung wurde wieder eine Mediationsanalyse mit dem PROCESS-Makro (Modell 4) für R von Hayes (2022) durchgeführt, wobei der Influencer-Typ (KI vs. Mensch) als unabhängige Variable (X), die parasoziale Beziehung als Mediator (M) und die Anzeigen-Einstellung als abhängige Variable (Y) im Modell definiert wurden. Die Analyse ergab, dass der Influencer-Typ einen signifikanten positiven Einfluss auf die parasoziale Beziehung hatte (Pfad a: $b = 0,845$; $t(234) = 4,72$; $p < 0,001$). Die parasoziale Beziehung wiederum hatte einen signifikanten positiven Einfluss auf die Anzeigen-Einstellung (Pfad b: $b = 0,407$; $t(233) = 7,91$; $p < 0,001$). Der direkte Effekt des Influencer-Typs auf die Anzeigen-Einstellung war nicht signifikant (Pfad c': $b = 0,060$; $t(233) = 0,41$; $p = 0,683$), der totale Effekt war hingegen signifikant (Pfad c: $b = 0,404$ ($p = 0,008$)). Der indirekte Effekt über die parasoziale Beziehung war ebenfalls signifikant ($b = 0,344$; $[0,191; 0,527]$). Damit wird bestätigt, dass der Einfluss des Influencer-Typs auf die Anzeigen-Einstellung nicht nur über die Kongruenz, sondern auch über die parasoziale Beziehung vollständig vermittelt wird. Zusätzlich bestätigte eine multiple Regression ($R^2 = 0,4$) mit Kongruenz, parasozialer Beziehung und Produkt-Involvement als Prädiktoren für die Anzeigen-Einstellung, dass sowohl Kongruenz ($\beta = 0,317$, $p < 0,001$) als auch die parasoziale Beziehung ($\beta = 0,233$, $p < 0,001$) signifikant zur Anzeigen-Einstellung beitragen, wobei Kongruenz den stärkeren Effekt aufweist. Auch das Produkt-Involvement hatte einen signifikanten Einfluss ($\beta = 0,114$, $p = 0,001$). Die Annahmen der Normalverteilung und Homoskedastizität der Residuen sowie das Fehlen von Multikollinearität wurden erfolgreich geprüft.

H8: Es wurde mittels separater Regressionen untersucht, ob die Produktkategorie den Effekt der parasozialen Beziehung auf die Anzeigen-Einstellung moderiert, also ob der positive Einfluss der parasozialen Beziehung in der Technik-Kategorie im

Vergleich zur Kosmetik-Kategorie abgeschwächt wird. In der Technik-Kategorie zeigte die Regression einen signifikanten Effekt der parasozialen Beziehung ($\beta_t = 0,334$; $t(114) = 4,942$; $p < 0,001$) und einen signifikanten Einfluss des Produkt-Involvements ($\beta = 0,181$; $t(114) = 3,104$; $p < 0,001$). In der Kosmetik-Kategorie war der Effekt der parasozialen Beziehung ebenfalls signifikant ($\beta_k = 0,411$; $t(116) = 5,618$; $p < 0,001$), während der Einfluss des Produkt-Involvements nicht signifikant war ($p = 0,095$). Die Annahmen der Normalverteilung und Homoskedastizität der Residuen wurden erfolgreich geprüft, sowie das Fehlen der Multikollinearität. Die Ergebnisse zeigen, dass der Einfluss der parasozialen Beziehung auf die Anzeigen-Einstellung in der Technik-Kategorie geringer ist ($\beta_t = 0,334$ vs. $\beta_k = 0,411$). Da jedoch nicht geprüft wurde, ob dieser Unterschied statistisch signifikant ist, konnte H8 nicht bestätigt werden.

H9: Es wurde untersucht, ob eine höhere Anzeigen-Einstellung zu einer höheren Marken-Einstellung (H9a), einer höheren Marken-Folgeabsicht (H9b) und einer höheren Engagement-Rate (H9c) führt. Zur Überprüfung dieser Hypothese wurden separate lineare Regressionsmodelle berechnet. Bei H9a zeigte das Modell einen signifikant positiven Effekt der Anzeigen-Einstellung auf die Marken-Einstellung ($\beta = 0,281$; $t(233) = 5,377$; $p < 0,001$) und einen signifikanten Einfluss des Produkt-Involvements ($\beta = 0,344$; $t(233) = 10,219$; $p < 0,001$), somit wird H9a bestätigt. Da die Residuen der Marken-Folgeabsicht nicht normalverteilt waren, wurde die Variable logarithmiert. Das Regressionsmodell zeigte einen signifikant positiven Effekt der Anzeigen-Einstellung ($\beta = 0,167$; $t(233) = 6,410$; $p < 0,001$) und einen signifikanten Effekt des Produkt-Involvements ($\beta = 0,09$; $t(233) = 5,287$; $p < 0,001$). H9b wurde somit bestätigt. H9c wurde mittels eines ordinalen Regressionsmodells analysiert. Es zeigte sich ein signifikanter positiver Effekt der Anzeigen-Einstellung ($\beta = 0,886$; $z = 4,652$; $p < 0,001$). Das bedeutet, dass eine Erhöhung der Anzeigen-Einstellung um eine Einheit die Wahrscheinlichkeit, einen höheren Engagement-Score (z. B. 1 = Like, Kommentar oder Teilen) zu erreichen, um das 2,43-Fache ($e^{0,886} = 2,43$) erhöht. Es zeigte sich ebenso ein signifikanter Einfluss des Produkt-Involvements ($\beta = 0,434$; $z = 2,741$; $p = 0,006$). Somit wurde auch H9c bestätigt. Die Annahmen der Normalverteilung und Homoskedastizität der Residuen wurden erfolgreich geprüft, sowie das Fehlen der Multikollinearität. Insgesamt wurden H9a, H9b und H9c somit bestätigt.

5. Diskussion

Conclusio und Interpretation

Diese Arbeit untersucht, inwiefern sich KI-Influencer von menschlichen Influencern hinsichtlich der Werbewirksamkeit auf Instagram unterscheiden und welche Rolle die Produktkategorie dabei spielt. Die Ergebnisse liefern dazu ein klares Bild: Menschliche Influencer erzielen in den Bereichen Anzeigen-Einstellung, Kongruenz und parasoziale Beziehungen signifikant höhere Werte als KI-Influencer. In Bezug auf Marken-Einstellung und Marken-Folgeabsicht zeigten sich hingegen keine statistisch signifikanten Unterschiede zwischen den beiden Influencer-Typen. Bei der Engagement-Rate erzielten menschliche Influencer zwar deutlich höhere Werte, jedoch war die Anzahl der Interaktionen zu gering, um einen statistisch signifikanten Effekt nachzuweisen. Die Mediationsanalysen bestätigen, dass sowohl Kongruenz als auch parasoziale Beziehungen den Einfluss des Influencer-Typs auf die Anzeigen-Einstellung vollständig vermitteln. Demnach erzielen menschliche Influencer eine höhere Anzeigen-Einstellung, weil sie als kongruenter zum Produkt wahrgenommen werden und stärkere parasoziale Beziehungen aufbauen. Die Regressionsanalysen zeigen, dass die Anzeigen-Einstellung wiederum einen statistisch signifikant positiven Einfluss auf die Marken-Einstellung, die Marken-Folgeabsicht und die Engagement-Rate hat. In der Technik-Kategorie erzielten beide Influencer-Typen bessere Ergebnisse als in der Kosmetik-Kategorie, mit Ausnahme von zwei Fällen, in denen die Kosmetik-Kategorie minimal besser abschnitt. Obwohl in der Technik-Kategorie insgesamt höhere Werte erzielt wurden, traten keine statistisch signifikanten Interaktionseffekte auf, sodass keine Moderation durch die Produktkategorie festgestellt wurde. Die höheren Werte in der Technik-Kategorie lassen sich darauf zurückführen, dass die Marken-Einstellung und das Produkt-Involvement bei Spotify statistisch signifikant höher ausfielen.

Daraus lässt sich als Antwort auf die Forschungsfrage ableiten, dass menschliche Influencer hinsichtlich der Anzeigen-Einstellung, Kongruenz und parasozialen Beziehungen eine höhere Wirkung erzielen als KI-Influencer. In Bezug auf die Marken-Einstellung, die Marken-Folgeabsicht (und die Engagement-Rate) zeigten sich vergleichbare Werte zwischen beiden Influencer-Typen. Diese Effekte traten dabei unabhängig vom beworbenen Produkt bzw. der Produktkategorie auf.

Die Ergebnisse stehen teilweise im Einklang mit früheren Studien und meinen Hypothesen, weichen aber in einigen Punkten davon ab:

Dass sich keine statistisch signifikanten Unterschiede in der Marken-Einstellung zwischen den Influencer-Typen zeigten, entspricht den Ergebnissen von Thomas und Fowler (2021). Niedrigere parasoziale Beziehungen von KI-Influencern wurden bereits in Interviews von Lou et al. (2023) und Koles et al. (2024) gezeigt. Auch positive Effekte der parasozialen Beziehungen auf die Werbeeffektivität zeigten frühere Studien (Chung & Cho, 2017; Yuan et al., 2019). Ebenso bestätigt sich, dass menschliche Influencer bei der Anzeigen-Einstellung besser abschneiden. Zwar fand Franke et al. (2023) keinen direkten Effekt, jedoch wurden dennoch höhere Werte für menschliche Influencer festgestellt.

Unerwartet zeigte sich jedoch, entgegen meiner Hypothese, ein statistisch signifikanter Unterschied in der Kongruenz zwischen den Influencer-Typen, wobei der menschliche Influencer besser abschnitt. Auch bei Franke et al. (2023) wurden höhere Kongruenzwerte für menschliche Influencer festgestellt, jedoch ohne statistische Signifikanz. Zudem konnte die Produktkategorie entgegen den Erwartungen keine statistisch signifikante moderierende Wirkung entfalten, obwohl eine entsprechende Tendenz vorhanden war. Dass KI-Influencer in der Technik-Kategorie nicht statistisch signifikant besser abschnitten, obwohl sie thematisch besser zu Technik- als zu Kosmetikprodukten passen sollten, könnte darauf zurückzuführen sein, dass die Probanden den KI-Influencern und deren Werbeanzeigen eher skeptisch gegenüberstanden. In den offenen Kommentaren äußerten sich Teilnehmende negativ gegenüber KI-Influencern, da ihnen das Menschliche fehle, sie diese generell ablehnten, die Werbung als wenig glaubwürdig empfanden und bezweifelten, dass KI-Influencer echte Produkterfahrungen vermitteln können. Frühere Studien bestätigen ebenfalls eine geringere Akzeptanz und Authentizität von KI-Influencern (Sands et al., 2022; Lou et al., 2023; Koles et al., 2024; Song et al., 2024). Dadurch konnte sich der Matchup-Effekt nach Kamins (1990) möglicherweise nicht entfalten, da die allgemeine Skepsis gegenüber KI-Influencern überwog und ein potenzieller Fit zu technischen Produkten in der Wahrnehmung der Probanden gar nicht zur Wirkung kam.

Die entgegen meiner Hypothese niedrigere Engagement-Rate und Marken-Folgeabsicht könnte auf mehrere Faktoren zurückzuführen sein. Erstens bauen KI-Influencer schwächere parasoziale Beziehungen auf. Zweitens wird ihre Kongruenz mit den beworbenen Produkten als geringer wahrgenommen. Sowohl die niedrigere Kongruenz als auch die niedrigeren parasozialen Beziehungen wirken sich negativ auf die Anzeigen-Einstellung aus und senken damit sowohl die Engagement-Rate als auch die Marken-Folgeabsicht. Drittens spiegeln die kritischen Kommentare der Probanden eine skeptische Haltung gegenüber KI-Influencern wider. Viertens könnte sich der anfängliche Neuheitseffekt abgenutzt haben, sodass das durch KI-Influencer geweckte Interesse im Laufe der Zeit nachließ. Diese Faktoren, also die schwächeren parasozialen Beziehungen, die geringere Kongruenz, die daraus resultierende niedrigere Anzeigen-Einstellung, die Skepsis der Probanden sowie der verblassende Neuheitseffekt, könnten gemeinsam erklären, warum die Engagement-Rate und die Marken-Folgeabsicht niedriger als erwartet ausfielen.

Theoretische und praktische Implikationen

Die Ergebnisse dieser Studie liefern neue Erkenntnisse zur Wirkung von KI-Influencern im Vergleich zu menschlichen Influencern und erweitern den bisherigen Forschungsstand in mehreren zentralen Punkten. Nach meinem Wissensstand wurden die Marken-Folgeabsicht und die Engagement-Rate erstmals bei KI-Influencern untersucht, wodurch eine bestehende Forschungslücke geschlossen wird. Ebenso wurde die parasoziale Beziehung erstmals quantitativ untersucht, wodurch erstmalig ein kausaler Nachweis erbracht wurde, dass die parasoziale Beziehung bei KI-Influencern signifikant schwächer ausgeprägt ist als bei menschlichen Influencern. Darüber hinaus wurden zwei neue Produktkategorien (Spotify und Labello) einbezogen, sodass frühere Erkenntnisse zu Kongruenz, Marken-Einstellung und Anzeigen-Einstellung unter neuen Bedingungen überprüft werden konnten. Zudem wurde der Einfluss parasozialer Beziehungen auf die Anzeigen-Einstellung in diesem Zusammenhang erstmals untersucht und gezeigt, dass sie neben der Kongruenz den Effekt des Influencer-Typs auf die Anzeigen-Einstellung vollständig mediiert. Darüber hinaus wurde ein Nachweis erbracht, dass die Anzeigen-Einstellung wiederum die Marken-Folgeabsicht, die Marken-Einstellung und die Engagement-Rate positiv beeinflusst. Schließlich wurde analysiert, wie sich die Wirkung von

parasozialen Beziehungen auf die Anzeigen-Einstellung je nach Produktkategorie unterscheidet, wobei sich ein schwächerer Effekt in der Technik-Kategorie abzeichnet.

Für Manager ist es wichtig zu wissen, dass menschliche Influencer in den untersuchten Aspekten besser abschneiden als KI-Influencer, auch wenn die Unterschiede teils gering sind. Besonders relevant ist, dass die parasozialen Beziehungen bei KI-Influencern schwächer ausgeprägt sind und die Kongruenz unabhängig vom Produkt niedriger ist und sich auch bei einem entsprechenden Fit nicht erhöht. Beides wirkt sich negativ auf die Anzeigen-Einstellung aus, die wiederum die Marken-Einstellung, die Marken-Folgeabsicht und die Engagement-Rate negativ beeinflusst. Um erfolgreiche Influencer-Kampagnen zu gestalten, sollten daher Influencer gewählt werden, die eine starke parasoziale Beziehung zu ihren Followern aufbauen und gut zur Marke passen, und dies gelingt aus aktueller Sicht menschlichen Influencern besser. Dennoch haben KI-Influencer potenzielle Vorteile. Sands et al. (2022) zeigen, dass sie eine höhere Word-of-Mouth erzeugen. Zudem werden ihre Anzeigen als neuartiger wahrgenommen, was die Innovationskraft der Marke steigern kann (Franke et al., 2023). Auch meine Studie sowie Thomas und Fowler (2021) zeigen, dass sich die Marken-Einstellung zwischen beiden Influencer-Typen kaum unterscheidet. Lou et al. (2023) stellten fest, dass KI-Influencer das Markenimage positiv beeinflussen können, da sie mit Innovation und Neuheit assoziiert werden.

Limitationen und zukünftige Forschung

Trotz der sorgfältigen Studiengestaltung weist diese Untersuchung einige Limitationen auf. Die begrenzte Repräsentativität der Stichprobe schränkt die Generalisierbarkeit der Ergebnisse ein. Zudem wurde keine echte KI-Influencerin verwendet, was zwar die interne Validität stärkte, jedoch die externe Validität beeinträchtigte. Ein weiterer Aspekt betrifft die Messung der parasozialen Beziehung, die nur in ihrer Anfangsphase erfasst werden konnte, da die Probanden die Influencer zuvor nicht kannten und keine wiederholte Exposition stattfand. Hinsichtlich der statistischen Analysen könnte ein umfassenderes Modell für H8, das die statistische Signifikanz überprüft, sowie eine parallele Mediationsanalyse mit Kongruenz und parasozialen Beziehungen und moderierende Mediationen zusätzliche Erkenntnisse liefern.

Zukünftige Forschung sollte mit Followern echter KI-Influencer und Followern menschlicher Influencer quasi-experimentelle Tests durchführen, um insbesondere die parasozialen Beziehungen weiter zu untersuchen. Darüber hinaus wäre es sinnvoll, weitere Produktkategorien einzubeziehen, um die Übertragbarkeit der Ergebnisse auf unterschiedliche Kontexte zu prüfen. Auch wenn vollständig autonome KI-Influencer derzeit noch nicht existieren, könnten sie in Zukunft relevant werden, weshalb ihre potenzielle Wirkung zukünftig untersucht werden sollte.

Einen innovativen Ansatz bieten KI-Influencer, die über Blockchain basierte Smart Contracts und Plattformen wie Virtual Protocol (Virtuals, 2025) gesteuert werden könnten. Im Gegensatz zu zentral gesteuerten Influencern könnten hier nicht nur Unternehmen oder Einzelpersonen, sondern auch die Community aktiv über die Parameter, wie Charaktereigenschaften und Werte, der KI-Influencer und deren Kooperationen und Vergütungsmodelle über die Blockchain mitentscheiden. Besonders relevant wäre in diesem Zusammenhang, ob die durch Blockchain-Technologie ermöglichte Transparenz, Fälschungssicherheit und Dezentralität das Vertrauen der Follower stärkt und welche Auswirkungen dies auf die Wahrnehmung von KI-Influencern hätte.

Mein persönliches Interesse an diesem Thema geht über die werberelevanten Erkenntnisse zu KI-Influencern hinaus. Besonders spannend finde ich, wie Menschen mit KI interagieren, welche Einflüsse sie auf unser Verhalten hat und welche gesellschaftlichen Veränderungen sie in Zukunft mit sich bringen könnte. Diese Überlegungen haben mich motiviert, mich intensiv mit dem Thema auseinanderzusetzen und durch meine Masterarbeit einen Beitrag zum besseren Verständnis unserer zunehmend durch KI geprägten Welt zu leisten.

Transparenzhinweis

Im Rahmen der Seminararbeit bei Herrn Univ.-Prof. Mag. Dr. Heribert Reisinger habe ich in einer Gruppenarbeit am Thema *KI in der Marketingkommunikation* gearbeitet und dabei eigenständig ein Kapitel zu KI-Influencern verfasst. Dieses Kapitel bildete den Ausgangspunkt für das Exposé dieser Masterarbeit. Es wurden keine Inhalte aus der Seminararbeit übernommen, sondern lediglich die in der Seminararbeit verwendeten Studien zu KI-Influencern (Thomas & Fowler, 2021; Sands et al., 2022; Franke et al., 2023; Lou et al., 2023) in die theoretischen Grundlagen dieser Arbeit einbezogen.

Die KI-Tools (DeepL Write und GPT4) wurden ausschließlich zur Lektorierung genutzt, vergleichbar mit einem menschlichen Lektorat, um die sprachliche Qualität zu verbessern und den Ausdruck zu optimieren.

Literaturverzeichnis

- Aw, E. C.-X., & Chuah, S. H.-W. (2021). “Stop the unattainable ideal for an ordinary me!” fostering parasocial relationships with social media influencers: The role of self-discrepancy. *Journal of Business Research*, 132, 146–157. <https://doi.org/10.1016/j.jbusres.2021.04.025>
- Bador, E. [@emilybador]. (2022, 9.September). Good :) morning :) [Instagram-Beitrag]. Instagram. <https://www.instagram.com/p/Ci95YlJN4ey/>
- Bador, E. [@emilybador]. (2025). Instagram-Profil. Instagram. https://www.instagram.com/darth_bador/?hl=de
- Belanche, D., Casaló, L. V., & Flavián, M. (2024). Human versus virtual influences, a comparative study. *Journal of Business Research*, 173, 114493. <https://doi.org/10.1016/j.jbusres.2023.114493>
- Breves, P., Amrehn, J., Heidenreich, A., Liebers, N., & Schramm, H. (2021). Blind trust? The importance and interplay of parasocial relationships and advertising disclosures in explaining influencers’ persuasive effects on their followers. *International Journal of Advertising*, 40(7), 1209–1229. <https://doi.org/10.1080/02650487.2021.1881237>
- Chung, S., & Cho, H. (2017). Fostering Parasocial Relationships with Celebrities on Social Media: Implications for Celebrity Endorsement. *Psychology & Marketing*, 34(4), 481–495. <https://doi.org/10.1002/mar.21001>
- Conde, R., & Casais, B. (2023). Micro, macro and mega-influencers on instagram: The power of persuasion via the parasocial relationship. *Journal of Business Research*, 158, 113708. <https://doi.org/10.1016/j.jbusres.2023.113708>

- Europäische Union. (2024). *Verordnung (EU) 2024/1689 des Europäischen Parlaments und des Rates vom 1. August 2024 zur Festlegung harmonisierter Vorschriften für Künstliche Intelligenz (AI Act) und zur Änderung bestimmter unionsrechtlicher Rechtsakte*. <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A32024R1689>
- Fleck, N., Korchia, M., & Le Roy, I. (2012). Celebrities in advertising: Looking for congruence or likeability? *Psychology & Marketing*, 29(9), 651–662. <https://doi.org/10.1002/mar.20551>
- Franke, C., Groeppel-Klein, A., & Müller, K. (2023). Consumers' responses to virtual influencers as advertising endorsers: Novel and effective or uncanny and deceiving? *Journal of Advertising*, 52(4), 523–539. <https://doi.org/10.1080/00913367.2022.2154721>
- Franke, C., & Groeppel-Klein, A. (2024). The role of psychological distance and construal level in explaining the effectiveness of human-like vs. cartoon-like virtual influencers. *Journal of Business Research*, 185, 114916. <https://doi.org/10.1016/j.jbusres.2024.114916>
- Figma. (2025). Instagram social post mockup (community). Figma. [https://www.figma.com/design/umBAtaRX0ayo1v0jxWmjds/Instagram-Social-Post-Mockup-\(Community\)?node-id=2-2&t=fg9M8lpQgxGG4CfK-0](https://www.figma.com/design/umBAtaRX0ayo1v0jxWmjds/Instagram-Social-Post-Mockup-(Community)?node-id=2-2&t=fg9M8lpQgxGG4CfK-0)
- Grand View Research. (2024). Virtual influencer market size, share & trends. Grand View Research. <https://www.grandviewresearch.com/industry-analysis/virtual-influencer-market-report>
- Google Cloud. (2024). *Was ist künstliche Intelligenz (KI)?* Google. <https://cloud.google.com/learn/what-is-artificial-intelligence>
- Gutuleac, R., Baima, G., Rizzo, C., & Bresciani, S. (2024). Will virtual influencers overcome the uncanny valley? The moderating role of social cues. *Psychology & Marketing*, 41(7), 1419–1431. <https://doi.org/10.1002/mar.21989>

- Hayes, A. F. (2018). Partial, conditional, and moderated moderated mediation: Quantification, inference, and interpretation. *Communication Monographs*, 85(1), 4–40. <https://doi.org/10.1080/03637751.2017.1352100>
- Holbrook, M. B., & Hirschman, E. C. (1982). The experiential aspects of consumption: Consumer fantasies, feelings, and fun. *Journal of Consumer Research*, 9(2), 132–140. <https://doi.org/10.1086/208906>
- Horton, D., & Richard Wohl, R. (1956). Mass Communication and Para-Social Interaction: Observations on Intimacy at a Distance. *Psychiatry (Washington, D.C.)*, 19(3), 215–229. <https://doi.org/10.1080/00332747.1956.11023049>
- Huang, M. H., & Rust, R. T. (2018). Artificial intelligence in service. *Journal of Service Research*, 21(2), 155–172. <https://doi.org/10.1177/1094670517752459>
- Kamins, M. A. (1990). An investigation into the “match-up” hypothesis in celebrity advertising: When beauty may be only skin deep. *Journal of Advertising*, 19(1), 4–13. <https://doi.org/10.1080/00913367.1990.10673175>
- Kim, I., Ki, C.-W., Lee, H., & Kim, Y.-K. (2024). Virtual influencer marketing: Evaluating the influence of virtual influencers’ form realism and behavioral realism on consumer ambivalence and marketing performance. *Journal of Business Research*, 176, 114611. <https://doi.org/10.1016/j.jbusres.2024.114611>
- Koles, B., Audrezet, A., Moulard, J. G., Ameen, N., & McKenna, B. (2024). The authentic virtual influencer: Authenticity manifestations in the metaverse. *Journal of Business Research*, 170, 114325. <https://doi.org/10.1016/j.jbusres.2023.114325>
- Kumar, V., B. Rajan, R. Venkatesan, and J. Lecinski. 2019. Understanding the role of artificial intelligence in personalized engagement marketing. *California Management Review* 61 (4):135–55. doi:10.1177/0008125619859317

- Lee, Y., & Koo, J. (2015). Athlete endorsement, attitudes, and purchase intention: The interaction effect between athlete endorser-product congruence and endorser credibility. *Journal of Sport Management*, 29, 523–538. <https://doi.org/10.1123/JSM.2014-0195>
- Leung, F., Gu, F., Li, Y., Zhang, J., & Palmatier, R. (2022). Influencer marketing effectiveness. *Journal of Marketing*, 86, 93–115. <https://doi.org/10.1177/00222429221102889>
- Lou, C. (2022). Social Media Influencers and Followers: Theorization of a Trans-Parasocial Relation and Explication of Its Implications for Influencer Advertising. *Journal of Advertising*, 51(1), 4–21. <https://doi.org/10.1080/00913367.2021.1880345>
- Lou, C., Kiew, S. T. J., Chen, T., Lee, T. Y. M., Ong, J. E. C., & Phua, Z. (2023). Authentically fake? How consumers respond to the influence of virtual influencers. *Journal of Advertising*, 52(4), 540–557. <https://doi.org/10.1080/00913367.2022.2149641>
- Lopez, A. [@fit_aitana]. (2024, 27. Oktober). *En el marco de la*. [Instagram-Beitrag]. Instagram. https://www.instagram.com/p/DBo69VgCASP/?img_index=1
- MacKenzie, S., Lutz, R., & Belch, G. (1986). The role of attitude toward the ad as a mediator of advertising effectiveness: A test of competing explanations. *Journal of Marketing Research*, 23, 130–143. <https://doi.org/10.1177/002224378602300205>
- Mrad, M., Ramadan, Z., Toth, Z., Nasr, L., & Karimi, S. (2024). Virtual Influencers Versus Real Connections: Exploring the Phenomenon of Virtual Influencers. *Journal of Advertising*, 1–19. <https://doi.org/10.1080/00913367.2024.2393711>
- Mori, M., MacDorman, K., & Kageki, N. (2012). The Uncanny Valley. *IEEE Robotics & Automation Magazine*, 19(2), 98–100. doi:10.1109/MRA.2012.2192811

- Muniz, F., Stewart, K., & Magalhães, L. (2024). Are they humans or are they robots? The effect of virtual influencer disclosure on brand trust. *Journal of Consumer Behaviour*, 23(3), 1234–1250. <https://doi.org/10.1002/cb.2271>
- Miquela, L. [@lilmiquela]. (2022, 26.September). Just a pixilated girl in this very real world. [Instagram-Beitrag]. Instagram. <https://www.instagram.com/p/Cxqok-LuG3f/>
- Miquela, L. [@lilmiquela]. (2024). Instagram-Profil. Instagram. <https://www.instagram.com/lilmiquela/?hl=de>
- Nass, C., & Moon, Y. (2000). Machines and mindlessness: Social responses to computers. *Journal of Social Issues*, 56(1), 81–103. <https://doi.org/10.1111/0022-4537.00153>
- Petty, R., & Cacioppo, J. (1984). Source factors and the elaboration likelihood model of persuasion. *Advances in Consumer Research*, 11, 668–672.
- Piwik. (2025). Engagement Rate. <https://piwikpro.de/glossar/engagement-rate/>
- PollPool. (2025). Find Survey Participants. PollPool. <https://www.poll-pool.com>
- R Core Team. (2020). R: A Language and Environment for Statistical Computing [Software]. R Foundation for Statistical Computing. <https://www.R-project.org/>
- Sands, S., Campbell, C. L., Plangger, K., & Ferraro, C. (2022). Unreal influence: Leveraging AI in influencer marketing. *European Journal of Marketing*, 56(6), 1721–1747. <https://doi.org/10.1108/EJM-12-2019-0949>
- Shaw, A. [@im_alice87]. (2024). Instagram-Profil. Instagram. https://www.instagram.com/im_alice87/

- Silvera, D. H., & Austad, B. (2004). Factors predicting the effectiveness of celebrity endorsement advertisements. *European Journal of Marketing*, 38(11/12), 1509–1526. <https://doi.org/10.1108/03090560410560218>
- Song, X., Lu, Y., & Yang, Q. (2024). The negative effect of virtual endorsers on brand authenticity and potential remedies. *Journal of Business Research*, 185, 114898. <https://doi.org/10.1016/j.jbusres.2024.114898>
- StataCorp. (2024). *Stata: Stata is statistical software for data science* [Software]. <https://www.stata.com/>
- Statista. (2024, 6.Mai). Ausgaben von Unternehmen für Influencer-Marketing auf Social Media weltweit bis 2024. *Statista*. <https://de.statista.com/statistik/daten/studie/1325021/umfrage/ausgaben-von-unternehmen-fuer-influencer-marketing-weltweit/>
- SoSci Survey. (2025). Online survey tool. SoSci Survey. <https://www.soscisurvey.de>
- SurveyCircle. (2025). Forschungswebseite SurveyCircle. SurveyCircle. <https://www.surveycircle.com>
- Tajfel, H., & Turner, J. C. (1986). The social identity theory of intergroup behavior. In S. Worchel & W. G. Austin (Eds.), *Psychology of intergroup relations* (pp. 7–24). Nelson-Hall.
- Thomas, V. L., & Fowler, K. (2021). Close encounters of the AI kind: Use of AI influencers as brand endorsers. *Journal of Advertising*, 50(1), 11–25. <https://doi.org/10.1080/00913367.2020.1810595>
- Till, B. D., & Busler, M. (2000). The match-up hypothesis: Physical attractiveness, expertise, and the role of fit on brand attitude, purchase intent and brand beliefs. *Journal of Advertising*, 29(3), 1–13. <https://doi.org/10.1080/00913367.2000.10673613>

- Trope, Y., & Liberman, N. (2010). Construal-level theory of psychological distance. *Psychological Review*, 117(2), 440. <https://doi.org/10.1037/a0018963>
- Virtuals. (2025). Agent Commerce Protocol. Virtuals. <https://app.virtuals.io/research/agent-commerce-protocol>
- Volles, B. K., Park, J., Van Kerckhove, A., & Geuens, M. (2024). How and when do virtual influencers positively affect consumer responses to endorsed brands? *Journal of Business Research*, 183, 114863. <https://doi.org/10.1016/j.jbusres.2024.114863>
- Yuan, C., Moon, H., Kim, K., & Wang, S. (2019). The influence of parasocial relationship in fashion web on customer equity. *Journal of Business Research*. <https://doi.org/10.1016/j.jbusres.2019.08.039>
- Zhou, X., Yan, X., & Jiang, Y. (2024). Making Sense? The Sensory-Specific Nature of Virtual Influencer Effectiveness. *Journal of Marketing*, 88(4), 84–106. <https://doi.org/10.1177/00222429231203699>