

www.iu.de

IU DISCUSSION

PAPERS

Business & Management

Künstlich, aber wirkungsvoll?

Eine empirische Analyse zur Kennzeichnungspflicht von KI-generierten Bildern und ihren Konsequenzen auf die Werbewirkung

FLORIAN PERST

NOEL SCHMITT

CONSTANTIN SCHUBART

IU Internationale Hochschule

Main Campus: Erfurt
Juri-Gagarin-Ring 152
99084 Erfurt

Telefon: +49 421.166985.23
Fax: +49 2224.9605.115
Kontakt/Contact: kerstin.janson@iu.org

Autorenkontakt/Contact to the author(s):

Dr. Florian Perst. ORCID-ID: 0009-0001-0426-6434 (Open Researcher und Contributor ID)
florian.perst@iu.org

Noel Schmitt ORCID-ID: 0009-0000-3692-8474 (Open Researcher und Contributor ID)
noel.schmitt@gmx.de

Prof. Dr. Constantin Schubart ORCID-ID: 0009-0008-9259-0533 (Open Researcher und Contributor ID)
constantin.schubart@iu.org

IU Discussion Papers, Reihe: Business & Management, Vol. 7, No. 16 (AUG 2026)

ISSN: 2750-0683

DOI: <https://doi.org/10.56250/4145>

Website: <https://repository.iu.org>

KÜNSTLICH, ABER WIRKUNGSVOLL?

EINE EMPIRISCHE ANALYSE ZUR KENNZEICHNUNGSPFLICHT VON KI-GENERIERTEN BILDERN UND IHREN KONSEQUENZEN AUF DIE WERBEWIRKUNG

**Florian Perst
Noel Schmitt
Constantin Schubart**

ABSTRACT:

This study examined the effects of the labelling requirement for AI-generated image content, introduced by Regulation (EU) 2024/1689, on advertising effectiveness in commercial tourism marketing. The aim was to investigate whether the legally mandated transparency disclosure in an advertisement negatively influences key advertising effectiveness dimensions. The variables examined included recall, perceived credibility, attitude toward the ad, trust, and purchase intention. Using a quantitative online survey with an experimental A/B design (N = 144), consumer responses to an advertisement featuring a real photograph were compared with those to a visually equivalent advertisement labelled as AI-generated. The results revealed no significant differences between the groups, indicating that the transparency disclosure did not reduce advertising effectiveness for the visually high-quality stimulus used here. Companies may therefore be able to leverage the efficiency potential of generative AI without necessarily suffering losses in advertising effectiveness. However, the absence of significant differences does not constitute evidence of equivalence, and the transferability of these findings to other industries, motifs, and target groups requires further empirical investigation.

KEYWORDS:

artificial intelligence, AI labelling, advertising effectiveness, transparency requirement, deepfake, AI regulation, Regulation (EU) 2024/1689

JEL classification: M31, M37, D91, O33, K24

AUTOR:INNEN



Dr. Florian Perst ist Dozent an der IU Internationalen Hochschule im Fachbereich Business & Management mit den Schwerpunkten Marketing, Personal- und Prozessmanagement. Sein Forschungsschwerpunkt liegt im Bereich der digitalen Customer Journeys, wobei er Schlüsselfaktoren der Kundeninteraktion und -bindung entlang digitaler Journeys untersucht.



Noel Schmitt ist gelernter Mediengestalter und hat seinen Bachelor in Online Marketing absolviert. Im Rahmen seiner Abschlussarbeit hat er sich auf die Dimensionen der Werbewirkung von KI-generierten Bildern fokussiert. Heute verbindet er die Erkenntnisse und Perspektiven als Projektmanager in einer Design- und Werbeagentur, wo er die Entwicklung und Integration von KI-Lösungen in Kunden- und Kampagnenprozessen begleitet.



Dr. Constantin Schubart ist seit 2020 Professor für Allgemeine Betriebswirtschaftslehre an der IU Internationale Hochschule im Dualen Studium am Standort Erfurt. Seine Schwerpunkte liegen im Bereich Managerial Economics. Er verfügt über mehr als 20 Jahre Berufserfahrung in Lehre, Forschung und Praxis

Einleitung

Die Leistungsfähigkeit von Text-zu-Bild-Modellen nimmt rasch zu. Nahezu täglich werden neue oder weiterentwickelte Systeme vorgestellt, und aktuelle Benchmark-Analysen zeigen die hohe Bildqualität und Realitätsnähe der führenden Modelle (Artificial Analysis, 2025). Dieser Fortschritt verwischt die Wahrnehmungsgrenze zwischen fotografischer Abbildung und künstlicher Erzeugung. Der Marketingbereich hat diese Entwicklung adaptiert: 98 % der befragten Agenturen setzen KI bereits ein, investieren in entsprechende Technologien und entwickeln teilweise eigene KI-Modelle (Observatory International Hamburg & BVDW, 2025, S. 7, 10, 37); die Bildgenerierung zählte bereits 2024 zu den als besonders relevant eingestuften Anwendungsfeldern (Statista, 2024b). Gleichzeitig äußern 60 % der befragten Marketingfachkräfte die Sorge, ihre berufliche Rolle könnte durch KI gefährdet sein – ein deutlicher Anstieg gegenüber 35 % im Vorjahr (Influencer Marketing Hub, 2024, zit. n. Statista, 2025). Die größte Herausforderung beim Einsatz von KI sehen 43 % in der Wahrung der Authentizität der Inhalte, während bereits 20 % die Kennzeichnungspflicht selbst als zentrale Hürde betrachten (Capterra, 2024, zit. n. Statista, 2024a).

Auf Konsumentenseite fällt die Bewertung uneinheitlich aus. Eine Befragung von YouGov (2025) zeigt, dass rund die Hälfte der Befragten Unbehagen gegenüber KI-generierten Social-Media-Posts empfindet; bei Artikeln und Blogbeiträgen sind es 44 %, bei Videos 41 % (YouGov, 2025, S. 16). Dabei stehen Millennials (27 %) und die Generation Z (26 %) KI-generierten Bildern positiver gegenüber als die Generation X (14 %) und die Baby-Boomer-Generation (10 %) (YouGov, 2025, S. 12). Komplementäre Ergebnisse liefert eine Umfrage zur Kampagne der Fashionmarke „Mango“ mit KI-generierten Modellen: 54,3 % der Befragten gaben an, ihre Loyalität gegenüber der Marke sei ungebrochen, während 81 % eine Kennzeichnungspflicht für KI-generierte Inhalte fordern (Appinio, 2024).

Mit der Verordnung (EU) 2024/1689 (KI-VO) wird eine solche Kennzeichnungspflicht verbindlich. Für werbetreibende Unternehmen entsteht daraus ein strategisches Dilemma zwischen zwei gegenläufigen Handlungsoptionen: Wer KI-generierte Bildmotive einsetzt, realisiert Effizienzgewinne in Produktionszeit und -kosten, muss den künstlichen Ursprung jedoch offenlegen, den Erstellungsprozess dokumentieren und mögliche negative Konsumentenreaktionen auf den Hinweis in Kauf nehmen. Wer weiterhin auf konventionell produzierte Fotografie setzt, vermeidet diese Risiken, verzichtet aber auf die Effizienzpotenziale. Welche Option vorzuziehen ist, hängt maßgeblich davon ab, wie stark der Transparenzhinweis die Werbewirkung tatsächlich beeinflusst.

Bisherige experimentelle Studien deuten auf negative Effekte des KI-Einsatzes im werblichen Umfeld hin: Belanche et al. (2025) zeigten reduzierte Nutzungs- und Empfehlungsabsichten bei KI-Bildern im Gastgewerbe, insbesondere bei hedonisch geprägten und als persönlich wichtig eingestuften Angeboten; Baek et al. (2024) wiesen eine verringerte Anzeigenglaubwürdigkeit und Spendenbereitschaft im prosozialen Kontext nach. Großskalige Felddaten von Exner et al. (2025) zeigen demgegenüber, dass KI-generierte Anzeigenmotive bei den Klickraten mit realen Motiven mithalten – allerdings vor allem dann, wenn ihre Künstlichkeit für Konsumenten nicht erkennbar ist (Exner et al., 2025, S. 22, 25). Genau diese Voraussetzung entfällt unter einer verpflichtenden Kennzeichnung.

Das Ziel des vorliegenden Papers ist es, die Auswirkungen der gesetzlich vorgeschriebenen Kennzeichnung von KI-generierten Bildern auf die Werbewirkung kommerzieller Anzeigen zu quantifizieren. Im Mittelpunkt steht die Frage, wie sich der Hinweis auf den künstlichen Ursprung des

Bildmotivs auf zentrale Dimensionen der Werbewirkung aus Sicht der Konsumenten auswirkt. Die zentrale Forschungsfrage hierfür lautet:

Wie verändert die gesetzlich vorgeschriebene Kennzeichnung eines KI-generierten Bildes im kommerziellen Kontext die Werbewirkung einer Anzeige im Vergleich zu einer Anzeige mit einem realen Foto?

Untersucht wird, ob sich durch die Kennzeichnung Unterschiede in der Erinnerungsleistung, in der wahrgenommenen Glaubwürdigkeit und im Vertrauen gegenüber der Anzeige sowie in der Einstellung zur Anzeige und in der Kaufintention ergeben. Der Untersuchungsfokus liegt auf der unmittelbaren Wirkung der Kennzeichnung im Kontext statischer, kommerzieller Anzeigen. Kanal- und branchenübergreifende Generalisierungen, Langzeiteffekte sowie alternative Kennzeichnungsformen werden nicht betrachtet.

Bevor die theoretischen Grundlagen zur Werbewirkung erläutert werden, ist zunächst zu klären, wann ein KI-generiertes Bild rechtlich als Deepfake gilt und somit der Kennzeichnungspflicht unterliegt.

Theoretischer Hintergrund

BEGRIFFSDEFINITION „DEEPPFAKE“ NACH DER VERORDNUNG (EU) 2024/1689

Der Begriff „Deepfake“ ist ein Kofferwort, das sich aus den englischen Begriffen „Deep Learning“ und „fake“ zusammensetzt (Baek et al., 2024, S. 15). Der Unionsgesetzgeber definiert einen Deepfake als „einen durch KI erzeugten oder manipulierten Bild-, Ton- oder Videoinhalt, der wirklichen Personen, Gegenständen, Orten, Einrichtungen oder Ereignissen ähnelt und einer Person fälschlicherweise als echt oder wahrheitsgemäß erscheinen würde“ (Verordnung (EU) 2024/1689, Art. 3 Nr. 60). Entscheidend ist dabei das Täuschungspotenzial: Die Verordnung rückt die Perspektive des Betrachters in den Vordergrund, wobei die bloße Eignung zur Täuschung ausreicht. Ein Nachweis, dass tatsächlich jemand getäuscht wurde, ist nicht erforderlich (Martini & Wendehorst, 2024, S. 265). Der Maßstab hierfür ist außerdem niedrig anzulegen, da weite Teile der Bevölkerung nur über geringe Kompetenzen zur Erkennung künstlich erzeugter Inhalte verfügen (Martini & Wendehorst, 2024, S. 265). Ein Indiz für das begrenzte Problembewusstsein liefert eine repräsentative Umfrage, nach der ein großer Teil der Bevölkerung den Begriff Deepfake nicht kennt oder dessen Bedeutung nicht erklären kann (Bitkom e. V., 2024). Aus der fehlenden Begriffskennntnis lässt sich zwar nicht unmittelbar auf eine geringe Erkennungsfähigkeit schließen; sie deutet jedoch darauf hin, dass die Möglichkeit künstlich erzeugter Inhalte von vielen Betrachtern nicht aktiv mitgedacht wird.

Ein KI-generiertes Werbebild erfüllt somit alle Tatbestandsmerkmale: Es ist KI-generiert, es ähnelt in der Regel der Kategorie wirklicher Personen, Gegenstände oder Orte, und es ist darauf ausgelegt, einer Person fälschlicherweise als echt zu erscheinen, unabhängig davon, ob es eine reale oder fiktive Szene darstellt.

Die daraus resultierende Kennzeichnungspflicht wirft die Frage auf, wie Konsumenten diesen Transparenzhinweis kognitiv verarbeiten und welche Auswirkungen dies auf ihre Wahrnehmung der Werbeanzeige hat.

GRUNDLAGENMODELLE

Um die Wirkung werblicher Reize und einzelner Elemente eines Werbemittels (wie der Kennzeichnung KI-generierter Bilder) auf Konsumenten zu analysieren, ist ein Verständnis der zugrunde liegenden kognitiven Prozesse innerhalb des Konsumenten notwendig. Dieser Abschnitt bildet die theoretische Grundlage, indem es zwei etablierte Modelle der Konsumentenverhaltensforschung miteinander in Beziehung setzt. Als Ausgangspunkt dient das S-O-R-Modell, das den grundlegenden Rahmen für die Verarbeitung externer Reize im Organismus liefert. Aufbauend auf den vorangegangenen Ausführungen wird im Folgenden das Persuasion Knowledge Model (PKM) von Friestad und Wright (1994) näher beschrieben. Das Modell thematisiert die Rolle des erlernten Wissens von Konsumenten bezüglich Überzeugungstaktiken und fand auch in der Studie von Baek et al. (2024) Anwendung.

S-O-R-MODELL

In der Konsumentenverhaltensforschung hat sich das S-O-R-Modell nach Mehrabian und Russel (1974) als grundlegendes Paradigma etabliert, um die Wirkung von Werbemaßnahmen auf Konsumenten systematisch zu analysieren (Hochreiter et al., 2023, S. 8; Schwarz & Hutter, 2012, S. 49). Dieses Modell

stellt eine Weiterentwicklung der einfacheren Stimulus-Response-Modelle (S-R-Modelle) dar, die aus der behavioristischen Psychologie stammen (Schwarz & Hutter, 2012, S. 49). S-R-Modelle gehen davon aus, dass ein externer Reiz (Stimulus), wie beispielsweise eine Werbeanzeige, unmittelbar eine beobachtbare Reaktion (Response) beim Konsumenten auslöst, beispielsweise einen Kauf (Rashedi, 2024, S. 23). Die internen psychischen Vorgänge, die zwischen Reiz und Reaktion ablaufen, werden dabei jedoch nicht berücksichtigt und das Individuum wird als eine sogenannte „Black Box“ betrachtet (Schwarz & Hutter, 2012, S. 49). Das auf dem Neobehaviorismus basierende S-O-R-Modell erweitert diesen Ansatz, indem es die im Organismus (O) ablaufenden, nicht direkt beobachtbaren Prozesse in die Analyse einbezieht (Hochreiter et al., 2023, S. 8). Ein von außen einwirkender Stimulus löst demnach im Inneren des Konsumenten psychische Vorgänge aus, die wiederum die schlussendliche Reaktion beeinflussen (Schwarz & Hutter, 2012, S. 49). Diese Variablen lassen sich in zwei Prozessarten unterteilen (Schwarz & Hutter, 2012, S. 50): Aktivierende Prozesse gelten als die zentralen Antriebskräfte des menschlichen Verhaltens und umfassen Emotionen, Motivation und Einstellungen. Sie sind für die emotionale Aufladung und die grundsätzliche Handlungsbereitschaft verantwortlich (Schwarz & Hutter, 2012, S. 50; Weinberg, 2003). Kognitive Prozesse dienen hingegen der gedanklichen Steuerung des Verhaltens. Hierzu zählen vor allem die Informationssuche, die Informationsverarbeitung sowie Lernprozesse (Schwarz & Hutter, 2012, S. 51). Aktivierende und kognitive Prozesse bedingen sich dabei gegenseitig und laufen nicht isoliert voneinander ab (Schwarz & Hutter, 2012, S. 50).

Übertragen auf den Marketingkontext fungiert eine Werbeanzeige als Stimulus (S). Die darin enthaltenen Elemente, wie beispielsweise eine explizite Kennzeichnung eines KI-generierten Bildes, werden im Organismus (O) des Betrachters verarbeitet. Diese internen Prozesse beeinflussen die Reaktion (R), die sich in der wahrgenommenen Glaubwürdigkeit, dem Vertrauen, der Einstellung zur Werbung oder der Kaufabsicht äußern kann. Das S-O-R-Modell liefert hierfür den grundlegenden Rahmen, lässt jedoch offen, wie genau Konsumenten werbliche Stimuli mental verarbeiten, insbesondere wenn sie diese als Überzeugungsversuch erkennen. Um diese spezifischen kognitiven Vorgänge im Organismus zu beleuchten, wird im Folgenden das Persuasion Knowledge Model herangezogen.

PERSUASION KNOWLEDGE MODEL

Im Marketingkontext interagieren zwei Akteure: der Werbetreibende („Agent“) und der Konsument („Target“). Diese Interaktion findet in einer sogenannten „Persuasion Episode“ statt, die den wahrgenommenen Überzeugungsversuch und die Bewältigungsstrategien des Konsumenten umfasst (Friestad & Wright, 1994, S. 2–3). Um diesen Überzeugungsversuch zu verarbeiten, greift der Konsument auf drei zentrale Wissensstrukturen zurück: das „Topic Knowledge“ (Wissen über das Produkt), das „Agent Knowledge“ (Einstellungen gegenüber dem Werbetreibenden) und das „Persuasion Knowledge“ (Verständnis für den Überzeugungsprozess selbst) (Friestad & Wright, 1994, S. 2–3). Diese drei Wissensstrukturen stehen in einem dynamischen Zusammenspiel und werden je nach Situation unterschiedlich aktiviert (Friestad & Wright, 1994, S. 3–4).

Der zentrale psychologische Mechanismus ist das „Change-of-Meaning“-Prinzip: der Prozess, bei dem eine zuvor neutrale Handlung als mögliche, bewusste Überzeugungstaktik erkannt und neu bewertet wird (Friestad & Wright, 1994, S. 12; Rahmani, 2023, S. 15). Sobald ein Konsument eine Taktik erkennt,

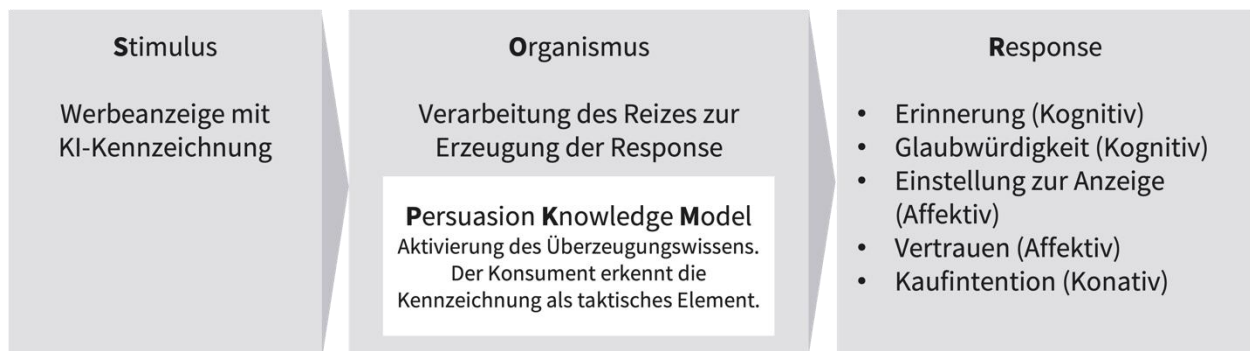
kann ein „Ablösungseffekt“ eintreten – die Zielperson löst sich mental von der Werbebotschaft und beginnt, die Motive des Werbetreibenden kritisch zu hinterfragen (Friestad & Wright, 1994, S. 13). Die Aktivierung des Überzeugungswissens führt zu einer Neubewertung, die je nach Taktik sowohl positiv als auch negativ ausfallen kann (Friestad & Wright, 1994, S. 13).

Das PKM erweist sich daher als geeignetes Instrument zur Analyse der Effekte von Transparenzhinweisen. Die Kennzeichnung eines KI-generierten Bildes kann den Prozess der Taktikerkennung auslösen und somit das Überzeugungswissen der Konsumenten aktivieren (Friestad & Wright, 1994, S. 13). Ob diese Aktivierung zu abwehrenden oder akzeptierenden Reaktionen führt, hängt von der individuellen Bewertung der wahrgenommenen Taktik ab.

DIMENSIONEN DER WERBEWIRKUNG

Die Werbewirkung lässt sich in drei Ebenen unterteilen: die kognitive, affektive und konative Dimension (Kühn, 2022, S. 236). Auf der kognitiven Ebene stehen die Werbeerinnerung als Maß für die Verankerung im Gedächtnis sowie die Glaubwürdigkeit als subjektive Zuschreibung der Wahrhaftigkeit einer Botschaft im Vordergrund (Bentele & Seidenglanz, 2008, S. 346; Naderer & Matthes, 2014, S. 1). Die affektive Ebene umfasst die Einstellung zur Anzeige („Attitude toward the Ad“), die als Prädiktor für den Werbeerfolg gilt, sowie das Vertrauen in das beworbene Angebot (Bentele & Seidenglanz, 2008, S. 346; Brown & Stayman, 1992, S. 34). Die konative Ebene bildet die handlungsorientierte Stufe und wird durch die Kaufintention operationalisiert, die als vermittelnde Variable zwischen Einstellung und tatsächlichem Verhalten fungiert (Mirabi et al., 2015, S. 268; Morwitz et al., 2007, S. 348).

Abb. 1: S-O-R Modell erweitert um das PKM



Quelle: Eigene Darstellung

Methodik

Zur Beantwortung der Forschungsfrage wurde eine quantitative Online-Befragung mit integriertem experimentellem A/B-Test durchgeführt, da dieses Design die Identifikation von Ursache-Wirkungs-Beziehungen ermöglicht (Tomczak et al., 2023, S. 46). Auf Grundlage des PKM und des Authentizitätskonzepts wurden für alle fünf Wirkungsdimensionen gerichtete Hypothesen formuliert, die jeweils einen negativen Effekt der KI-Kennzeichnung annehmen. Die Teilnehmer wurden per Zufallsprinzip zwei Gruppen zugewiesen: Gruppe A (n = 76) erhielt eine Anzeige mit realem Foto, Gruppe B (n = 68) eine visuell nahezu identische Anzeige mit dem Hinweis „KI-Generiertes Bild“. Die Erfassung der abhängigen Variablen erfolgte mittels etablierter Messinstrumente: Die Erinnerungsleistung wurde über Free Recall (offene Frage) und Aided Recall (geschlossene Fragen mit Distraktoren) erhoben (Brosius et al., 2014, S. 26). Glaubwürdigkeit und Einstellung zur Anzeige wurden jeweils über siebenstufige semantische Differenziale operationalisiert (MacKenzie & Lutz, 1989), das Vertrauen mittels einer adaptierten ADTRUST-Skala (Soh et al., 2009) und die Kaufintention über die elfstufige Juster-Skala (Juster, 1966, S. 672). Ergänzend wurden demografische Kontrollvariablen erhoben.

Die interne Konsistenz der mehrteiligen Skalen wurde über Cronbachs Alpha geprüft. Die Skalen zur Einstellung zur Anzeige ($\alpha = 0,78$) und zum Vertrauen ($\alpha = 0,83$) erreichten gute Werte. Die ursprünglich dreiteilige Glaubwürdigkeitsskala wies dagegen eine unzureichende Konsistenz auf ($\alpha = 0,48$). Nach Ausschluss des Adjektivpaars „unvoreingenommen/voreingenommen“, das von den Teilnehmern offenkundig uneinheitlich interpretiert wurde, ergab sich für die verbleibenden zwei Items ein Wert von $\alpha = 0,81$. Alle weiteren Auswertungen basieren auf dieser bereinigten Skala.

Als Stimuli dienten zwei Geschwisteranzeigen für eine fiktive Pauschalreise nach Kreta, die inhaltlich und gestalterisch identisch waren. Das Foto für Gruppe A stammte aus der Bilddatenbank Unsplash; für Gruppe B wurde daraus über die Plattform Freepik mithilfe des Modells FLUX.1 Kontext ein digitaler Zwilling generiert und unten rechts transparent mit dem Hinweis „KI-Generiertes Bild“ gekennzeichnet. Damit unterschieden sich die beiden Versionen ausschließlich in der unabhängigen Variablen dieser Untersuchung.

Abb. 2: Stimulus für Online-Umfrage Gruppe A

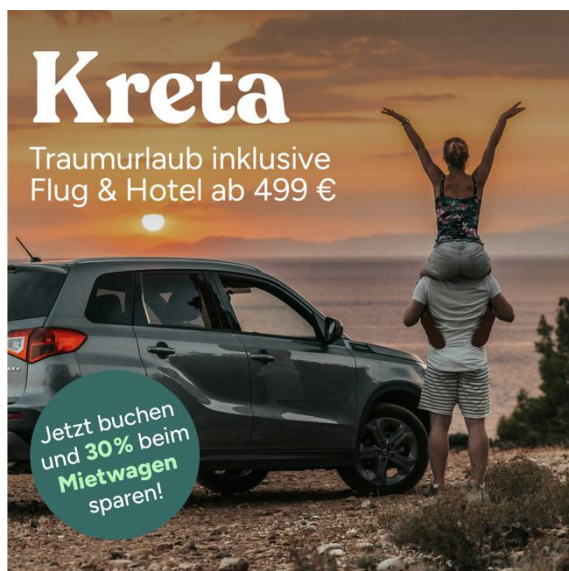


Abb. 3: Stimulus für Online-Umfrage Gruppe B



Quelle: Foto von Alex Meier, Unsplash.

Quelle: Eigene Generierung über Freepik mit
FLUX.1 Kontext.

Der Stimulus wurde für 10 Sekunden präsentiert, eine Dauer oberhalb der in der Literatur diskutierten Identifikationsschwelle (Elsen et al., 2025, S. 386). Die Zielgruppe umfasste in Deutschland lebende Instagram-Nutzer im Alter von 18 bis 34 Jahren; die Rekrutierung erfolgte über persönliche und universitäre Netzwerke. Um Erwartungseffekte zu vermeiden, wurde die Befragung unter dem neutralen Titel „Beurteilung von Werbeanzeigen im Tourismussektor“ ausgespielt. Nach Datenbereinigung verblieben $N = 144$ auswertbare Fälle, bestehend aus 91 Frauen, 52 Männern und einer diversen Person.

Ergebnisse

In diesem Abschnitt werden die empirischen Ergebnisse der Untersuchung dargestellt. Zunächst werden die Angaben zur Werbeerinnerung ausgewertet, einschließlich der offenen Antworten sowie der kontrollierenden Erinnerungsfragen. Anschließend folgt die deskriptive Darstellung der abhängigen Variablen beider Versuchsgruppen. Ergänzend werden innerhalb der Experimentalgruppe Unterschiede zwischen Teilnehmern mit und ohne explizite Erwähnung des KI-Hinweises beschrieben. Damit wird ein umfassender Überblick über die erhobenen Daten bereitgestellt, der die Grundlage für die anschließende inferenzstatistische Analyse bildet.

Die Häufigkeitsverteilung der kodierten Antworten auf die offene Frage zur Werbeerinnerung ist in Tabelle 1 dargestellt. Die meisten Teilnehmer gaben korrekte Antworten; irrelevante Antworten traten in Gruppe A etwas häufiger auf als in Gruppe B. Zur Quantifizierung der qualitativen Daten wurden die Antworten anhand eines vierstufigen Schemas kodiert. Als „Korrekt“ wurden Antworten gewertet, die ausschließlich richtige Aspekte und mindestens eine der drei zentralen Botschaften (Reiseziel Kreta, Preis 499 €, 30 % Rabatt) nannten. Antworten mit Detailfehlern (z. B. „Griechenland“ statt „Kreta“) fielen in die Kategorie „Teilweise Korrekt“, während Aussagen mit fundamentalen inhaltlichen Fehlern (z. B. falscher Preis) als „Falsch“ kodiert wurden. Fehlende oder themenfremde Eingaben wurden als „Irrelevant“ klassifiziert.

Tab. 1: Häufigkeitsverteilung der Antworten auf die offene Frage zur Werbeerinnerung nach Gruppe¹

	Korrekt	Teilw. Korrekt	Falsch	Irrelevant
Gruppe A	52	11	2	11
Gruppe B	50	10	2	6

Quelle: Eigene Darstellung.

Tabelle 2 zeigt, dass der Hinweis auf die KI-Kennzeichnung in den Freitextantworten zur Werbeerinnerung von der Mehrheit der Teilnehmer nicht explizit genannt wurde. Lediglich 25 % der Experimentalgruppe (17 von 68 Personen) erwähnten den Hinweis oder semantisch verwandte Begriffe.

¹ Anmerkung. $N = 144$. Dargestellt sind die absoluten Häufigkeiten der kodierten Freitextantworten pro Gruppe

Tab. 2: Häufigkeitsverteilung der in der Freitextantwort erwähnten KI-Kennzeichnung²

	Häufigkeit	Anteil
Kennzeichnung nicht genannt	51	75 %
Kennzeichnung genannt	17	25 %

Quelle: Eigene Darstellung.

Bei den Kontrollfragen zur Werbeerinnerung zeigten sich deskriptive Unterschiede zwischen den Gruppen. In Kontrollfrage 1 lag der Anteil falscher Antworten in Gruppe B bei 5,88 % gegenüber 1,32 % in Gruppe A. In Kontrollfrage 2 fiel der Unterschied stärker aus: In Gruppe B beantworteten 17,65 % der Teilnehmer die Frage falsch, in Gruppe A 7,89 %. Der Abstand zwischen den Gruppen beträgt damit 9,76 Prozentpunkte (siehe Tab. 3). Wie im folgenden Abschnitt gezeigt wird, sind diese Unterschiede statistisch nicht signifikant.

Tab. 3: Häufigkeitsverteilung der Antworten auf die Kontrollfragen zur Werbeerinnerung nach Gruppe³

	Kontrollfrage 1			Kontrollfrage 2		
	Korrekt	Falsch	Falsch (%)	Korrekt	Falsch	Falsch (%)
Gruppe A	75	1	1,32	70	6	7,89
Gruppe B	64	4	5,88	56	12	17,65

Quelle: Eigene Darstellung.

Die deskriptive Auswertung der abhängigen Variablen für beide Gruppen ist in Tabelle 4 und Tabelle 5 dargestellt. Insgesamt zeigen sich ähnliche Mittelwerte für die Skalen Einstellung zur Anzeige, Glaubwürdigkeit und Vertrauen in beiden Gruppen.

Tab. 4: Deskriptive Statistik der abhängigen Variablen für die Gruppe A⁴

	Mittelwert	Std. Abw.	Min	Max
Aad_Index	4,86	1,16	1,5	6,75
Glaubwuerdigkeit_Index	4,20	1,45	1,0	6,5
Vertrauen_Index	4,05	1,53	1,0	7,0
Kaufintention	2,68	2,26	0,0	9,0

Quelle: Eigene Darstellung.

Die Kaufintention ist in Gruppe B etwas höher ausgeprägt als in Gruppe A. Die Mittelwerte und Streuungen der abhängigen Variablen weisen in beiden Gruppen ähnliche Ausprägungen auf.

² Anmerkung. N = 68. Die Tabelle stellt die absoluten und prozentualen Häufigkeiten dar, mit denen der „KI-Generiertes Bild“-Hinweis oder semantisch verwandte Begriffe in den Freitextantworten zur Werbeerinnerung (Frage 1) explizit genannt wurden.

³ Anmerkung. N = 144 (Gruppe A N = 76, Gruppe B N = 68). Dargestellt sind die absoluten Häufigkeiten der korrekten und falschen Antworten sowie der prozentuale Anteil der falschen Antworten pro Gruppe. Kontrollfrage 1 bezog sich auf das Reiseziel (Kreta), Kontrollfrage 2 auf das Mietwagen-Angebot (30 % Preisnachlass).

⁴ Anmerkung. N = 76. Dargestellt sind die Mittelwerte, Standardabweichungen sowie Minimal- und Maximalwerte für die gebildeten Indizes. Alle Skalen reichten von 1 (niedrigste Ausprägung) bis 7 (höchste Ausprägung). Die Kaufintention wurde auf einer Skala von 0 (sehr gering) bis 10 (sehr hoch) erhoben.

Tab. 5: Deskriptive Statistik der abhängigen Variablen für die Gruppe B⁵

	Mittelwert	Std. Abw.	Min	Max
Aad_Index	4,83	1,18	2,0	7,0
Glaubwuerdigkeit_Index	4,34	1,38	1,67	7,0
Vertrauen_Index	4,04	1,39	1,67	7,0
Kaufintention	3,0	2,27	0,0	9,0

Quelle: Eigene Darstellung.

Im Unterschied dazu konzentriert sich die Kaufintention deutlich im unteren Wertebereich, was eine klar nicht-glockenförmige Verteilung erkennen lässt. Dieses Muster ist bei einem hypothetischen Angebot erwartbar: Kaufabsichten, die sich auf eine fiktive Anzeige ohne realen Buchungskontext beziehen, sind nur eingeschränkt aussagekräftig, da geäußerte Absicht und tatsächliches Verhalten insbesondere bei unbekannten Angeboten auseinanderfallen (Morwitz et al., 2007, S. 347). Die Kaufintention ist hier daher weniger als Prognose realen Kaufverhaltens denn als vergleichende Bewertungsgröße zwischen den beiden Experimentalgruppen zu interpretieren.

Zur weiterführenden Betrachtung wurde Gruppe B zusätzlich nach der expliziten Erwähnung des KI-Hinweises in der offenen Erinnerungsfrage differenziert. Für die Teilgruppe ohne Erwähnung (siehe Tab. 6) zeigen sich leicht erhöhte Mittelwerte in der Einstellung zur Anzeige sowie in der Kaufintention im Vergleich zur Teilgruppe mit Erwähnung des KI-Hinweises (siehe Tab. 7). Auch die wahrgenommene Glaubwürdigkeit liegt ohne Erwähnung leicht höher, während sich die Vertrauenswerte in beiden Teilgruppen auf einem ähnlichen Niveau befinden, mit einem geringfügigen Vorteil für die Gruppe ohne Erwähnung.

Tab. 6: Deskriptive Statistik der abhängigen Variablen für die Gruppe B „Keine KI-Erwähnung“⁶

	Mittelwert	Std. Abw.	Min	Max
Aad_Index	4,91	1,15	2,25	7,0
Glaubwuerdigkeit_Index	4,39	1,36	2,0	7,0
Vertrauen_Index	3,96	1,36	1,67	7,0
Kaufintention	3,35	2,29	0,0	9,0

Quelle: Eigene Darstellung.

Insgesamt weisen die deskriptiven Kennwerte damit auf leicht höhere Index-Werte bei den Teilnehmern hin, die den KI-Hinweis nicht aktiv in der Freitextantwort wiederholen, während bei aktiver Erwähnung ein leicht niedrigeres Bewertungsniveau vorliegt. In dieser weiterführenden Betrachtung wird auf eine inferenzstatistische Prüfung der Unterschiede zwischen den Teilgruppen verzichtet, da diese Analyse nicht Bestandteil des zuvor definierten Untersuchungsdesigns ist. Der Zweck dieser Auswertung liegt darin, innerhalb der Gruppe B eine ergänzende deskriptive Orientierung zu ermöglichen, ohne zusätzliche Hypothesen zu generieren oder über den geplanten Analyseumfang

⁵ Anmerkung. $N = 68$. Dargestellt sind die Mittelwerte, Standardabweichungen sowie Minimal- und Maximalwerte für die gebildeten Indizes. Alle Skalen reichten von 1 (niedrigste Ausprägung) bis 7 (höchste Ausprägung). Die Kaufintention wurde auf einer Skala von 0 (sehr gering) bis 10 (sehr hoch) erhoben.

⁶ Anmerkung. $N = 51$. Dargestellt sind die Mittelwerte, Standardabweichungen sowie Minimal- und Maximalwerte für die gebildeten Indizes. Alle Skalen reichten von 1 (niedrigste Ausprägung) bis 7 (höchste Ausprägung). Die Kaufintention wurde auf einer Skala von 0 (sehr gering) bis 10 (sehr hoch) erhoben.

hinauszugehen. Die Darstellung dient damit ausschließlich der Einordnung und Beschreibung der beobachtbaren Muster, während die inferenzstatistische Prüfung auf die im Forschungsdesign vorgesehenen Gruppenvergleiche beschränkt bleibt.

Tab. 7: Deskriptive Statistik der abhängigen Variablen für die Gruppe B „KI-Erwähnung“⁷

	Mittelwert	Std. Abw.	Min	Max
Aad_Index	4,57	1,26	2,0	6,0
Glaubwuerdigkeit_Index	4,18	1,46	1,5	6,0
Vertrauen_Index	4,27	1,49	2,33	6,67
Kaufintention	1,94	1,89	0,0	6,0

Quelle: Eigene Darstellung.

INFERENZSTATISTISCHE HYPOTHESENPRÜFUNG

Aufbauend auf der deskriptiven Analyse wird in diesem Abschnitt geprüft, ob die beobachteten Tendenzen und Unterschiede zwischen Gruppe A und Gruppe B statistisch signifikant sind. Je nach Skalenniveau der abhängigen Variablen kommen unterschiedliche inferenzstatistische Tests zum Einsatz. Für die intervallskalierten Indizes (Glaubwürdigkeit, Einstellung zur Anzeige, Vertrauen und Kaufintention) wird der *t*-Test für unabhängige Stichproben verwendet. Die Erinnerungsleistung wird anhand ordinaler Daten mit dem Mann-Whitney-U-Test überprüft, für die nominalskalierten Kontrollfragen kommt der Chi-Quadrat-Test zum Einsatz. Das Signifikanzniveau für die Ablehnung der Nullhypothesen wird auf das konventionelle Alpha-Niveau von $p < .05$ festgelegt.

Die Ergebnisse zeigen, dass für keine der abhängigen Variablen ein statistisch signifikanter Unterschied zwischen den Gruppen nachgewiesen werden konnte. Weder die freie Erinnerungsleistung ($U = 2426$, $Z = -0,79$, $p = .43$, $r = -0,07$) noch die Glaubwürdigkeit ($t(142) = -0,60$, $p = .55$), die Einstellung zur Anzeige ($t(142) = 0,18$, $p = .86$), das Vertrauen ($t(142) = 0,04$, $p = .97$) oder die Kaufintention ($t(142) = -0,84$, $p = .40$) unterschieden sich signifikant zwischen der Gruppe mit realem Foto und der Gruppe mit KI-Kennzeichnung. Auch die gestützte Erinnerungsleistung wies keine signifikanten Gruppenunterschiede auf (Kontrollfrage 1: $\chi^2(1, N = 144) = 2,23$, $p = .14$, $\phi = 0,12$; Kontrollfrage 2: $\chi^2(1, N = 144) = 3,12$, $p = .07$, $\phi = 0,15$). Die Effektstärken waren durchweg vernachlässigbar ($|d| \leq 0,14$); ergänzend gerechnete lineare Regressionen als Robustheitsprüfung erklärten praktisch keine Varianz ($R^2 \approx .00$). Sämtliche Nullhypothesen werden damit beibehalten. Das Ausbleiben signifikanter Unterschiede ist dabei nicht mit dem statistischen Nachweis von Gleichwertigkeit gleichzusetzen; ein solcher Nachweis würde ein Äquivalenztestverfahren und eine entsprechend geplante Fallzahl voraussetzen.

⁷ Anmerkung. $N = 17$. Dargestellt sind die Mittelwerte, Standardabweichungen sowie Minimal- und Maximalwerte für die gebildeten Indizes. Alle Skalen reichten von 1 (niedrigste Ausprägung) bis 7 (höchste Ausprägung). Die Kaufintention wurde auf einer Skala von 0 (sehr gering) bis 10 (sehr hoch) erhoben.

Diskussion

Die inferenzstatistischen Analysen ergaben für keine der geprüften abhängigen Variablen einen signifikanten Unterschied zwischen der Anzeige mit realem Foto und der Anzeige mit KI-Kennzeichnung. Dies gilt gleichermaßen für die kognitive Ebene (Erinnerungsleistung und Glaubwürdigkeit), die affektive Ebene (Einstellung zur Anzeige und Vertrauen) sowie die konative Ebene (Kaufintention). Der Befund legt nahe, dass die explizite Kennzeichnung eines KI-generierten Bildes unter den gegebenen Stimulus- und Kontextbedingungen keine messbare Beeinträchtigung zentraler Werbewirkungsmaße auslöst.

Aus Sicht des S-O-R-Modells sind die Befunde mit der Annahme vereinbar, dass der Stimulus „KI-Generiertes Bild“ im Organismus des Konsumenten keine ausgeprägte aktivierende oder kognitive Prozessveränderung ausgelöst hat, die sich in der anschließenden Befragung niedergeschlagen hätte. Dies zeigt sich zum einen an den nahezu identischen Mittelwerten der Gesamtgruppen und zum anderen an der Subdifferenzierung innerhalb der Gruppe B: Teilnehmer, die den Hinweis in der offenen Erinnerungsfrage nicht erwähnten, wiesen deskriptiv leicht höhere Mittelwerte für Einstellung zur Anzeige, Glaubwürdigkeit und Kaufintention auf, während die Vertrauenswerte beider Teilgruppen auf einem ähnlichen Niveau lagen. Da diese Teilgruppen nicht inferenzstatistisch geprüft wurden, sind die Unterschiede nicht als signifikant zu interpretieren; sie deuten jedoch darauf hin, dass die aktive Erinnerung des Hinweises allein keine systematische Abwertung der Gesamtbewertung begründet.

Nach der Logik des PKM könnte dies auf eine ausbleibende oder nur schwache Taktikerkennung seitens der Konsumenten zurückzuführen sein. Die Kennzeichnung wurde teils wahrgenommen und aktiv wiedergegeben, löste jedoch keinen „Change-of-Meaning“ im Sinne von Friestad und Wright (1994) aus, der die Anzeige als Überzeugungsversuch des Werbetreibenden neu gerahmt hätte (Friestad & Wright, 1994, S. 12–13). Auch bei den Teilnehmern, die die Kennzeichnung in ihrer freien Erinnerung aktiv wiedergaben, ließ sich kein nennenswerter Unterschied in der Werbewirkung feststellen. Das Muster der nahezu parallelen Bewertungen ist mit der Annahme vereinbar, dass Persuasion Knowledge nicht in relevantem Umfang aktiviert wurde, sodass der Hinweis kein eigenständiges Qualitätssignal etablierte und die Motive des Werbetreibenden nicht neu bewertet wurden.

Auch die Anschlusskonzepte Authentizität und Spillover liefern im vorliegenden Ergebnisbild eher Negativbefunde: Wenn Authentizität als Basis für Glaubwürdigkeit und Vertrauen fungiert, hätte eine KI-Kennzeichnung als potenzieller Bruch der Echtheitszuschreibung die Bewertungen senken müssen. Dass dies ausblieb, deutet darauf hin, dass die Kennzeichnung in den vorliegenden Daten nicht als starkes Authentizitätsargument gewertet wurde. Ebenso war ein Inter-Entitäts-Spillover von der Künstlichkeit des Bildes auf die Glaubwürdigkeit des Angebots theoretisch zu erwarten; er blieb in den vorliegenden Daten jedoch ohne Ausprägung.

MÖGLICHE ERKLÄRUNGEN UND ABGRENZUNGEN

Die ausbleibenden Effekte widersprechen der zugrunde liegenden Theorie nicht; ihre Ursachen lassen sich aus den vorliegenden Daten jedoch nicht direkt ableiten. Denkbar ist etwa, dass die Salienz des Hinweises hinsichtlich Gestaltung und Platzierung zu gering war, um eine Taktikerkennung im Sinne von Friestad und Wright (1994) auszulösen.

Einen weiterführenden Erklärungsansatz liefert die differenzierte Betrachtung der Erinnerungsleistung. Obwohl der Hinweis visuell wahrnehmbar platziert war, wurde er in der offenen Erinnerung nur von 25 % der Teilnehmer der Experimentalgruppe aktiv genannt. Dieser Befund deutet weniger auf bloße Unaufmerksamkeit hin als auf eine selektive Informationsverarbeitung. Vor dem Hintergrund des Limited Capacity Models (Lang, 2000, S. 56) werden eingehende Reize nach ihrer Relevanz gefiltert, bevor sie enkodiert und weiterverarbeitet werden; dieser Filterprozess läuft weitgehend automatisch und nicht als bewusste Entscheidung des Konsumenten ab. Dass 75 % der Teilnehmer den Hinweis potenziell sahen, ihn aber nicht in der freien Erinnerung wiedergaben, legt nahe, dass die Information „KI-generiertes Bild“ in diesem Kontext nur peripher enkodiert und als nicht entscheidungsrelevante Zusatzinformation verarbeitet wurde. Anders als Warnhinweise, die als Gefahrensignal eine hohe kognitive Aktivierung auslösen, scheint die KI-Kennzeichnung bei einem ästhetisch hochwertigen Bildmotiv diesen Relevanzfilter nicht zu passieren. Im Sinne des PKM fehlte damit der Auslöser für das kritische Überzeugungswissen: Die Kennzeichnung wurde nicht als Bruch der Authentizität gewertet, sondern als neutrale Eigenschaft akzeptiert oder übergangen.

Im Vergleich zur Studie von Baek et al. (2024) zur Offenlegung des künstlichen Ursprungs war die Kennzeichnung in der durchgeführten Studie sehr viel reduzierter und weniger raumeinnehmend. Die Kennzeichnung könnte also neben der fehlenden Relevanz auch als nebensächlich wahrgenommen worden sein oder nicht den emotionalen Aktivierungsschwellenwert wie in der Studie von Baek et al. (2024) erreicht haben.

Außerdem ist der Kontext zu berücksichtigen: In prosozialen Kampagnen wird die Offenlegung des künstlichen Ursprungs stärker emotional gerahmt und kann dadurch eine niedrigere Aktivierungsschwelle überschreiten als in kommerziellen Anzeigen. Auch zeigt die Studie konträre Ergebnisse zu Belanche et al. (2025). Dort schnitten KI-Bilder gegenüber realen Bildern schlechter ab, obwohl beide Untersuchungen im Bereich der Gastgewerbeprestationen angesiedelt sind. Dafür kommen mehrere Gründe in Betracht, die sich aus Motivwahl, Bildgestaltung und Verarbeitung ableiten lassen. In der Studie von Belanche et al. (2025) dienen Abbildungen von Hotelgebäuden ohne weitere textliche Beschreibung innerhalb des Bildes als Stimulus. In dieser Studie sahen die Teilnehmer eine vollwertige Werbeanzeige mit Überschrift, Textblock, Störerelement und Hintergrundbild. Die Anzeige zeigte außerdem Menschen in einem emotionalen Kontext. Solche „Stimmungsbilder“ können gefühlsvermittelnd wirken und werden durch die Textelemente gerahmt, was die Interpretation stärker lenkt.

Demgegenüber ermöglicht die isolierte Abbildung eines Hotelgebäudes eine detailliertere Echtheitsprüfung, die Unsicherheiten zur Übereinstimmung mit der Realität sichtbar macht. Bei der Generierung des Stimmungsbildes für die Anzeige wurde außerdem auf eine natürliche Ästhetik und einen möglichst „Nicht-KI-Look“ geachtet und Aspekte nach Exner et al. (2025) wie stark gesättigte Farben (Exner et al., 2025, S. 23) und überproportional große Gesichter (Exner et al., 2025, S. 23) wurden bewusst vermieden. Ergänzend ist möglich, dass Konsumenten in der Zeit zwischen den durchgeführten externen Studien und der vorliegenden Erhebung bereits differenziertes Überzeugungswissen aufgebaut haben und der Hinweis daher möglicherweise nicht als Qualitätsmangel gewertet wurde. Der vermehrte Einsatz von KI-generierten Inhalten in der Werbung könnte bereits zu einer gestiegenen Akzeptanz der Konsumenten geführt haben.

KRITISCHE REFLEXION

Ein methodischer Aspekt betrifft die Stichprobengröße. Die anvisierte Gesamtteilnehmerzahl von 385 Teilnehmern konnte nicht erreicht werden. Nach der Bereinigung standen lediglich 144 vollständige Datensätze zur Verfügung. Eine geringere Stichprobe reduziert grundsätzlich die Möglichkeit, sehr kleine Effekte statistisch zuverlässig nachzuweisen und schränkt damit die Präzision der Schätzungen ein. Zugleich zeigen die vorliegenden Ergebnisse keine Anhaltspunkte für ausgeprägte oder substanzielle Unterschiede zwischen den Gruppen. Vor diesem Hintergrund ist es zwar möglich, dass sich bei einer größeren Stichprobe geringfügige Effekte statistisch signifikant gezeigt hätten. Es erscheint jedoch unwahrscheinlich, dass sich ein grundlegend anderes Befundmuster ergeben hätte. Für zukünftige Untersuchungen wäre dennoch eine breiter angelegte Rekrutierungsstrategie sinnvoll, um die Aussagekraft insbesondere hinsichtlich kleiner Effekte zu erhöhen.

Hinzu kommen messbezogene Einschränkungen. Die Glaubwürdigkeitsskala musste nachträglich um ein Item gekürzt werden; dies erhöht die Messgenauigkeit, verringert jedoch die inhaltliche Breite des Konstrukts. Das Involvement wurde lediglich indirekt über den zeitlichen Abstand zum letzten Urlaub erfasst. Dieser Proxy zeigte keinen Zusammenhang mit der Kaufintention und dürfte die tatsächliche situative Involviertheit nur unzureichend abbilden. Da das Involvement maßgeblich steuert, wie intensiv Konsumenten eine Werbebotschaft verarbeiten, bleibt offen, ob die ausbleibenden Effekte eher auf ein insgesamt niedriges Involvement oder auf die geringe Salienz des Hinweises zurückzuführen sind. Schließlich handelt es sich um eine nicht-probabilistische Gelegenheitsstichprobe, die über persönliche und universitäre Netzwerke rekrutiert wurde und überdurchschnittlich weiblich sowie formal hoch gebildet ist. Eine Verallgemeinerung über die untersuchte Altersgruppe und den untersuchten Kontext hinaus ist daher nicht zulässig.

Über die Stichprobengröße hinaus betrifft eine weitere wesentliche Limitation die spezifische Ausgestaltung des eingesetzten Stimulus. Die Schlussfolgerungen dieser Untersuchung basieren auf einem einzelnen Stimulus, der sowohl hinsichtlich des Bildmotivs als auch der spezifischen Kennzeichnungsvariante festgelegt war. Dadurch ließ sich aus den Ergebnissen nicht ableiten, dass KI-generierte Bilder generell dieselbe oder eine ähnliche Werbewirkung entfalten können wie nicht-KI-generierte Bilder. Die Generalisierbarkeit der Befunde auf andere Branchen, visuelle Motive, Produktkategorien oder alternative Kennzeichnungsformen ist daher eingeschränkt. Wie die Studien von Baek et al. (2024) und Belanche et al. (2025) zeigen, können sowohl der inhaltliche Kontext als auch das visuelle Erscheinungsbild eines Mediums einen entscheidenden Einfluss auf die Wahrnehmung haben. Der Fokus der vorliegenden Studie lag auf dem Tourismusmarketing, wo Bildmotive häufig als „Stimmungsbilder“ fungieren, um eine Atmosphäre zu transportieren. Dies stellt einen anderen Anwendungsfall dar als beispielsweise im E-Commerce für Fast Moving Consumer Goods oder Fashion. Dort liegt der Effizienzhebel oft in der digitalen Platzierung realer Produkte in synthetischen Umgebungen, was logistisch aufwendige Shootings ersetzt. Die kritische Echtheitsprüfung durch den Konsumenten könnte hier unterschiedlich ausgeprägt sein: Während bei einem Urlaubsmotiv die echte Erfahrung im Vordergrund steht, ist bei physischen Produkten die authentische Wiedergabe von Material und Beschaffenheit entscheidend. Auch wenn externe Studien wie zur Mango-Kampagne (Appinio, 2024) eine hohe Akzeptanz bei KI-Models im Fashion-Bereich andeuten, lässt sich dies nicht pauschal übertragen. Die Toleranz gegenüber der Kennzeichnung könnte sinken, je stärker das KI-

Element die eigentlich zu bewertende Produkteigenschaft (z. B. Passform von Kleidung) direkt betrifft oder verfälscht. Um fundierte Aussagen über unterschiedliche Anwendungsbereiche treffen zu können, wären zukünftige Forschungsvorhaben notwendig, die verschiedene Branchen, Bildmotive und Kennzeichnungsvarianten in den Vergleich einbeziehen. Für weiterführende Forschung ergibt sich hieraus die Frage, ob in Branchen mit höherer Produktbezogenheit (etwa Fashion, FMCG, Kosmetik, Elektronik oder Immobilien) die Kennzeichnung stärkere Effekte auslöst. Zukünftige Studien sollten daher gezielt untersuchen, unter welchen Bedingungen die Taktikerkennung im Sinne des PKM aktiviert wird und wann die Echtheitsprüfung zu einer kritischeren Bewertung führt.

Außerdem können die Gestaltung, Wortwahl und Platzierung des Hinweises die Wahrnehmung stark prägen. Im Beispiel von Baek et al. (2024) wurde in einem Anzeigenmotiv der Hinweistext in fetter roter Schrift hervorgehoben. In der vorliegenden Studie war der Hinweis lediglich subtil am unteren Bildrand mit einem Symbol platziert. Auch Motivkategorie, Bildästhetik und persönliches Interesse am Thema spielen eine Rolle.

HANDLUNGSEMPFEHLUNGEN

Basierend auf den Ergebnissen dieser Arbeit können Unternehmen die Einführung der Kennzeichnungspflicht gemäß der KI-VO strategisch neu bewerten. Die zentrale Befürchtung, dass allein der Hinweis auf einen künstlichen Ursprung die Werbewirkung oder das Vertrauen signifikant beeinträchtigt, ließ sich unter den Bedingungen dieser Studie nicht bestätigen. Marketingentscheider sollten die Effizienzpotenziale von KI in der Bildproduktion daher prüfen, statt aus Sorge vor der Kennzeichnungspflicht grundsätzlich darauf zu verzichten. Die Daten deuten darauf hin, dass Konsumenten dem Label zumindest bei hochwertigen Bildmotiven und in einem vergleichbaren Kontext weitgehend indifferent gegenüberstehen. Der strategische Fokus sollte folglich weniger auf der Vermeidung der Kennzeichnung als auf der Qualität der KI-Integration liegen.

Erste Marktbeobachtungen zeigen, dass Unternehmen die Kennzeichnungspflicht bereits umsetzen. Die Berliner Volksbank veröffentlichte im August 2025 eine Kampagne für ihr Produkt „Girokonto blauorange“, deren Bild- und Videoinhalte vollständig KI-generiert sind; sowohl die Anzeigen auf Meta als auch die Motive der Kampagnenwebsite sind mit einem Hinweis auf den künstlichen Ursprung versehen (campaigngermany.de, 2025; meedia.de, 2025; Berliner Volksbank eG, 2025; Meta, 2025).

Über die Frage, ob KI eingesetzt wird, hinaus stellt die Gestaltung des Hinweises selbst eine strategische Variable dar. Theoretisch lässt sich annehmen, dass die Wortwahl die Aktivierungsschwelle des Überzeugungswissens beeinflusst: Formulierungen, die eine bloße Bearbeitung realen Ausgangsmaterials nahelegen, dürften weniger stark mit vollständiger Künstlichkeit assoziiert werden als ein Hinweis auf eine vollständige Generierung. Diese Annahme ist bislang empirisch nicht geprüft; zudem ist jede Formulierung an den Vorgaben der KI-VO zu messen, die eine klare und deutliche Offenlegung des künstlichen Ursprungs verlangt (Verordnung (EU) 2024/1689, Erwägungsgrund 134).

Für Unternehmen ergeben sich daraus zwei komplementäre Gestaltungsoptionen: Erstens können Motive bewusst so konzipiert werden, dass sie erkennbar als kreative Inszenierung angelegt sind und damit keine Echtheitsprüfung provozieren. Zweitens kann, sofern rechtlich zulässig, die Formulierung des Hinweises so gewählt werden, dass sie den Grad der KI-Beteiligung differenziert kommuniziert. Beide Ansätze zielen darauf ab, die Kennzeichnung als neutrale Information zu rahmen, anstatt sie zum

Auslöser kritischer Bewertungsprozesse werden zu lassen; ihre Wirksamkeit ist bislang jedoch nicht empirisch belegt.

Ein möglicher, in dieser Studie allerdings nicht variiertes Einflussfaktor ist die visuelle Qualität der generierten Bilder. Der eingesetzte Stimulus war technisch hochwertig und wies keine typischen Artefakte auf. Da nur diese eine Qualitätsstufe geprüft wurde, lässt sich aus den Daten kein Zusammenhang zwischen Bildqualität und Akzeptanz des Labels ableiten. Als Erklärungshypothese ist ein solcher Zusammenhang jedoch plausibel, da negative Effekte in externen Studien vor allem dort berichtet werden, wo Künstlichkeit als Qualitätsmangel wahrgenommen wird (Belanche et al., 2025). Zugleich ist zu berücksichtigen, dass sich die Bildqualität führender Text-zu-Bild-Modelle rasch verbessert hat (Artificial Analysis, 2025) und offensichtliche Artefakte daher seltener werden; der praktische Schwerpunkt verschiebt sich damit von der Vermeidung technischer Fehler hin zur gestalterischen Passung des Motivs. Ein definierter Qualitätsstandard, der etwa „Uncanny Valley“-Effekte (Kim et al., 2024), unnatürliche Proportionen oder einen stark gesättigten, überperfekten „KI-Look“ ausschließt, bleibt für Unternehmen und Agenturen dennoch sinnvoll.

Das Ziel dieser Maßnahmen sollte sein, negative Qualitätszuschreibungen durch den Konsumenten zu vermeiden, sodass der Transparenzhinweis eine formale Information bleibt und nicht als Warnsignal für Minderwertigkeit interpretiert wird. Trotz des insgesamt unauffälligen Befundmusters ist der individuelle Kontext bei der Ausspielung der Werbemittel zu berücksichtigen. Da Angehörige älterer Altersgruppen tendenziell skeptischer reagieren (YouGov, 2025, S. 12) und externe Studien im prosozialen Bereich negative Effekte fanden (Baek et al., 2024), ist ein pauschaler Ansatz risikobehaftet. Vor großflächigen Roll-Outs empfiehlt sich daher die Durchführung systematischer A/B-Tests in den relevanten Kanälen mit klar definierten Zielgruppen, um die Wirkung auf Klick- und Conversion-Raten im spezifischen Kontext zu validieren. Dies liefert eine belastbare Datengrundlage, um die erzielten Produktivitätsgewinne gegen potenzielle Wahrnehmungsrisiken abzuwägen. Um die ab August 2026 greifende Kennzeichnungspflicht rechtssicher umzusetzen, müssen zudem interne Prozesse frühzeitig angepasst werden. Unabhängig vom konkreten Einsatzszenario ist bereits jetzt eine lückenlose Dokumentation aller KI-generierten Inhalte zu etablieren. Unternehmen sollten diese Dokumentationspflicht vertraglich von beauftragten Agenturen einfordern und Mitarbeiter gezielt schulen. Dies schafft Rechtssicherheit und verhindert Compliance-Risiken, sobald die Übergangsfristen der KI-VO enden.

Fazit

Die vorliegende Arbeit widmete sich dem Spannungsfeld zwischen technologischer Effizienz und regulatorischer Transparenz in der Marketingkommunikation. Vor dem Hintergrund der Verordnung (EU) 2024/1689 wurde die Frage untersucht, ob die verpflichtende Kennzeichnung KI-generierter Bildinhalte als Warnsignal wirken und die Werbewirkung negativ beeinträchtigen kann. Die theoretische Herleitung über das PKM legte die Vermutung nahe, dass ein solcher Transparenzhinweis das Überzeugungswissen der Konsumenten aktivieren, die Wahrnehmung von Authentizität stören und somit zu abwehrenden Reaktionsmustern führen kann.

Diese theoretische Annahme konnte durch die empirische Untersuchung im Kontext einer fiktiven Tourismuskampagne nicht bestätigt werden. Der experimentelle Vergleich zwischen einer Anzeige mit realem Foto und einer visuell gleichwertigen, aber als künstlich gekennzeichneten Anzeige zeigte keine statistisch signifikanten Unterschiede auf den kognitiven, affektiven und konativen Wirkungsebenen. Weder die Glaubwürdigkeit der Werbebotschaft noch das Vertrauen in den Anbieter oder die Kaufintention wurden durch die Offenlegung des künstlichen Ursprungs gemindert. Der Befund leistet damit einen Beitrag zur Debatte um die Akzeptanz synthetischer Medien: Er spricht dagegen, dass die Kennzeichnung als solche zwangsläufig als Warnsignal wirkt. Welche Faktoren stattdessen über die Konsumentenreaktion entscheiden, lässt sich aus den Daten nicht bestimmen; die visuelle Qualität und Stimmigkeit des Materials ist hierfür eine naheliegende, in dieser Studie jedoch nicht variierte Erklärung. Solange KI-generierte Inhalte visuell überzeugen und keine offensichtlichen Artefakte aufweisen, dürfte der Transparenzhinweis eher als neutrale Zusatzinformation verarbeitet oder peripher übergangen werden, ohne einen negativen „Change of Meaning“ im Sinne eines Authentizitätsverlustes auszulösen.

Für zukünftige Forschung ergibt sich daraus die Notwendigkeit, die hier gewonnenen Befunde unter variierenden Bedingungen zu prüfen. Dazu gehören insbesondere Studien mit größeren Stichprobenumfängen, der Einsatz unterschiedlicher Branchen- und Motivkategorien, variierende Grade der Emotionalisierung und alternative Formen der Kennzeichnung. Auch die Untersuchung unterschiedlicher Bildqualitäten sowie stärker produkt- oder personenorientierter Motive erscheint vielversprechend, um die Moderatoren der Konsumentenreaktion auf KI-generierte Inhalte systematischer zu erfassen. Auf Basis der vorliegenden Ergebnisse ist davon auszugehen, dass die Werbewirkung nicht allein durch die Kennzeichnung bestimmt wird, sondern im Zusammenspiel aus Kontext, Motiv und visueller Qualität entsteht.

Auf Grundlage dieser Perspektive lassen sich nicht nur Forschungsbedarfe ableiten, sondern auch Implikationen für die praktische Anwendung erkennen. Für die Marketingpraxis liefern die Ergebnisse der vorliegenden Studie zwar keinen universellen Freifahrtschein, jedoch ein wichtiges Indiz für eine differenzierte Risikobewertung unter den neuen regulatorischen Rahmenbedingungen. Unternehmen sollten die technologischen Potenziale zur Kostensenkung und Beschleunigung der Inhaltsproduktion nicht pauschal aufgrund der Kennzeichnungspflicht verwerfen. Vielmehr empfiehlt es sich, den Einsatz von KI als potenziellen Wettbewerbsvorteil zu prüfen und in kontrollierten Testumgebungen zu validieren. Die Herausforderung der Zukunft liegt dabei weniger im Verstecken der künstlichen Intelligenz, sondern in der Sicherstellung einer Qualität, die den ästhetischen Erwartungen der Konsumenten standhält und somit die Akzeptanz des Transparenzhinweises begünstigt.

Abschließend erscheint plausibel, dass mit der flächendeckenden Einführung der Kennzeichnungspflicht ab August 2026 Gewöhnungseffekte eintreten. Ähnlich wie bei der Kennzeichnung von Produktplatzierungen könnte sich der KI-Hinweis als selbstverständlicher Transparenzstandard etablieren, der Orientierung schafft, ohne den werblichen Dialog nachhaltig zu stören. Ob dies eintritt, ist eine empirische Frage, die künftige Längsschnittstudien beantworten müssen.

Literaturverzeichnis:

- Appinio. (2024). *Appinio* | MKT: *Mango* KI *Kampagne*.
<https://research.appinio.com/#/de/survey/public/EiLHuXOcF>
- Artificial Analysis. (2025). *Text to Image Leaderboard* | *Artificial Analysis*. Artificial Analysis.
<https://artificialanalysis.ai/image/leaderboard/text-to-image>
- Baek, T. H., Kim, J., & Kim, J. H. (2024). Effect of disclosing AI-generated content on prosocial advertising evaluation. *International Journal of Advertising*, 1–22.
<https://doi.org/10.1080/02650487.2024.2401319>
- Belanche, D., Ibáñez-Sánchez, S., Jordán, P., & Matas, S. (2025). Customer reactions to generative AI vs. Real images in high-involvement and hedonic services. *International Journal of Information Management*, 85, 102954. <https://doi.org/10.1016/j.ijinfomgt.2025.102954>
- Bentele, G., & Seidenglanz, R. (2008). Vertrauen und Glaubwürdigkeit. In G. Bentele, R. Fröhlich, & P. Szyszka (Hrsg.), *Handbuch der Public Relations* (S. 346–361). VS Verlag für Sozialwissenschaften.
https://doi.org/10.1007/978-3-531-19667-1_25
- Berliner Volksbank eG. (2025). *Kostenloses Girokonto*. <https://www.berliner-volksbank.de/privatkunden/girokonto-karten/kostenloses-girokonto.html>
- Bitkom e. V. (2024, Mai 17). *Jeder und jede Dritte hat noch nie von „Deepfakes“ gehört* | *Presseinformation* | Bitkom e. V. <https://www.bitkom.org/Presse/Presseinformation/Jeder-und-jede-Dritte-hat-noch-nie-von-Deepfakes-gehoert>
- Brosius, H.-B., Fahr, A., & Kaut, V. (2014). *Werbewirkung im Fernsehen II: Befunde aus der Medienforschung*. Nomos Verlagsgesellschaft mbH & Co. KG.
<https://doi.org/10.5771/9783845253695>
- Brown, S. P., & Stayman, D. M. (1992). Antecedents and Consequences of Attitude Toward the Ad: A Meta-Analysis. *Journal of Consumer Research*, 19(1), 34. <https://doi.org/10.1086/209284>
- campaigngermany.de. (2025, 7. August). KI-Spots: So wirbt Denkwerk für das Junge-Leute-Angebot der Berliner Volksbank. Campaign. <https://campaigngermany.de/news/beitrag/1230-ki-spots-so-wirbt-denkwerk-fuer-das-junge-leute-angebot-der-berliner-volksbank.html>
- Elsen, M., Pieters, R., & Wedel, M. (2025). Effects of advertising exposure duration and frequency: A theory and initial test. *Journal of Marketing Analytics*, 13(2), 386–404.
<https://doi.org/10.1057/s41270-024-00373-4>
- Exner, Y., Hartmann, J., Netzer, O., Zhang, S., & Ding, Z. (2025). *AI in Disguise – Quasi-Experimental Analysis of a Large-Scale Deployment of AI-Generated Display Ads*. 32.
- Friestad, M., & Wright, P. (1994). The Persuasion Knowledge Model: How People Cope with Persuasion Attempts. *Journal of Consumer Research*, 21(1), 1–31. <https://doi.org/10.1086/209380>
- Hochreiter, V., Benedetto, C., & Loesch, M. (2023). The Stimulus-Organism-Response (S-O-R) Paradigm as a Guiding Principle in Environmental Psychology: Comparison of its Usage in Consumer Behavior and Organizational Culture and Leadership Theory. *Journal of Entrepreneurship and Business Development*, 3(1), 7–16. <https://doi.org/10.18775/jebd.31.5001>

- Juster, F. T. (1966). *Consumer Buying Intentions and Purchase Probability: An Experiment in Survey Design*. 658–696.
- Kim, W., Ryoo, Y., & Choi, Y. K. (2024). That uncanny valley of mind: When anthropomorphic AI agents disrupt personalized advertising. *International Journal of Advertising*, 1–30. <https://doi.org/10.1080/02650487.2024.2411669>
- Kühn, C. (2022). *Digitalisierung im stationären Reisevertrieb: Der Einfluss von In-Store-Technologien auf das Kaufverhalten*. Springer Fachmedien Wiesbaden. <https://doi.org/10.1007/978-3-658-37099-2>
- Lang, A. (2000). The Limited Capacity Model of Mediated Message Processing. *Journal of Communication*, 50(1), 46–70. <https://doi.org/10.1111/j.1460-2466.2000.tb02833.x>
- MacKenzie, S. B., & Lutz, R. J. (1989). An Empirical Examination of the Structural Antecedents of Attitude toward the Ad in an Advertising Pretesting Context. *Journal of Marketing*, 53(2), 48–65. <https://doi.org/10.1177/002224298905300204>
- Martini, M., & Wendehorst, C. (Hrsg.). (2024). *KI-VO: Verordnung über Künstliche Intelligenz: Kommentar*. C.H. Beck.
- meedia.de. (2025). Denkwerk-Kampagne: Berliner Volksbank setzt auf KI-Videos. MEEDIA. <https://meedia.de/news/beitrag/19856-berliner-volksbank-setzt-auf-ki-videos.html>
- Mehrabian, A., & Russell, J. A. (1974). *An approach to environmental psychology*. the MIT Press.
- Meta. (2025). Werbebibliothek. <https://www.facebook.com/ads/library/>
- Mirabi, D. V., Akbariyeh, H., & Tahmasebifard, H. (2015). *A Study of Factors Affecting on Customers Purchase Intention*. 2.
- Morwitz, V. G., Steckel, J. H., & Gupta, A. (2007). When do purchase intentions predict sales? *International Journal of Forecasting*, 23(3), 347–364. <https://doi.org/10.1016/j.ijforecast.2007.05.015>
- Naderer, B., & Matthes, J. (2014). Verfahren zur Messung der Werbewirkung und Werbeeffizienz. In F.-R. Esch, T. Langner, & M. Bruhn (Hrsg.), *Handbuch Controlling der Kommunikation* (S. 1–17). Springer Fachmedien Wiesbaden. https://doi.org/10.1007/978-3-658-05260-7_16-1
- Observatory International Hamburg & BVDW. (2025). *Treiber der Transformation: Wie Agenturen generative KI nutzen Ergebnisse der größten Agenturumfrage zur KI-Nutzung in Deutschland*. <https://www.bvdw.org/wp-content/uploads/2025/04/BVDW-x-Observatory-KI-Agenturbefragung.pdf>
- Rahmani, V. (2023). Persuasion knowledge framework: Toward a comprehensive model of consumers' persuasion knowledge. *AMS Review*, 13(1–2), 12–33. <https://doi.org/10.1007/s13162-023-00254-6>
- Rashedi, J. (2024). *Customer Insights: Kundenbedürfnisse und Konsumentenverhalten verstehen*. Springer Fachmedien Wiesbaden. <https://doi.org/10.1007/978-3-658-43392-5>
- Schwarz, U., & Hutter, K. (2012). Marketing-Management: Wie sich das Verhalten von Konsumenten beeinflussen lässt. In S. Hoffmann, U. Schwarz, & R. Mai (Hrsg.), *Angewandtes Gesundheitsmarketing* (S. 45–55). Springer Fachmedien Wiesbaden. https://doi.org/10.1007/978-3-8349-4035-3_4

- Soh, H., Reid, L. N., & King, K. W. (2009). Measuring Trust In Advertising. *Journal of Advertising*, 38(2), 83–104. <https://doi.org/10.2753/JOA0091-3367380206>
- Statista. (2024a). *Leading challenges of using generative artificial intelligence (GenAI) for social media marketing according to marketers worldwide as of May 2024*. Capterra; MarketingCharts. <https://www.statista.com/statistics/1490913/genai-challenges-social-media-marketing-worldwide/>
- Statista. (2024b, Januar). *Marketing areas where generative AI has greatest potential*. Statista. <https://www.statista.com/statistics/1460868/marketing-potential-gen-ai-worldwide/>
- Statista. (2025). *Share of marketers concerned artificial intelligence (AI) may jeopardize their roles worldwide in 2023 and 2024*. Influencer Marketing Hub. <https://www.statista.com/statistics/1607820/ai-concerns-marketers-jeopardize-role-worldwide/>
- Tomczak, T., Reinecke, S., & Gollnhofer, J. (2023). *Marketingplanung: Einführung in die marktorientierte Unternehmens- und Geschäftsfeldplanung*. Springer Fachmedien Wiesbaden. <https://doi.org/10.1007/978-3-658-34221-0>
- Verordnung (EU) 2024/1689 des Europäischen Parlaments und des Rates vom 13. Juni 2024 zur Festlegung harmonisierter Vorschriften für künstliche Intelligenz (Verordnung über künstliche Intelligenz).*
- Weinberg, P. (2003). *Konsumentenverhalten* (W. Kroeber-Riel, Hrsg.; 8., aktualisierte und ergänzte Auflage). Verlag Franz Vahlen.
- YouGov. (2025). *Trust or concern? How Germany feels about generative AI in media*. <https://platform.yougov.com/whitepapers?id=30215>