

MASTERARBEIT | MASTER'S THESIS

Titel | Title

AI-generated Advertising: The Impact of AI Disclosure and Involvement Levels on
Consumers' Purchase Decision and Brand Evaluation

verfasst von | submitted by

Jennifer Böhm

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:

Univ.-Prof. Dr. Martin Eisend M.A.

Abstract (English)

This study examines the effect of AI disclosure in advertisements on perceived authenticity, brand image, and purchase intention by differentiating between low and high involvement based on the Elaboration Likelihood Model by Petty and Cacioppo. A 2x2 between-subjects online experiment was conducted, manipulating AI disclosure and product involvement.

Based on prior research, it was expected that an AI label would negatively influence consumers evaluations, especially for high involvement products advertisements.

The results indicate that involvement does not act as a moderator, meaning that the level of involvement has no significant effect on consumer evaluations. However, regardless of involvement, AI disclosure increases authenticity and brand image. In addition, a significant mediation of authenticity on brand image as well as purchase intention is found.

These findings underscore the importance of how AI is communicated in advertising, suggesting that disclosure can promote perceptions of transparency and innovation.

Keywords: Advertising, AI, AI-generated ads, Artificial Intelligence, Authenticity, Brand Image, Elaboration Likelihood Model, Involvement, Marketing, Purchase Intention

Abstract (Deutsch)

Diese Studie untersucht die Auswirkungen der Kennzeichnung des Einsatzes von künstlicher Intelligenz (KI) in Werbung auf die wahrgenommene Authentizität, Kaufabsicht und das Markenbild, indem zwischen niedrigem und hohem Involvement auf der Grundlage des Elaboration Likelihood Models von Petty und Cacioppo unterschieden wird. Es wurde ein 2x2 Online-Experiment durchgeführt, bei dem die Offenlegung von KI und das Produktinvolvement manipuliert wurden.

Basierend auf vorheriger Recherche wurde erwartet, dass eine KI-Kennzeichnung die Bewertungen der Verbraucher negativ beeinflusst, insbesondere bei Werbung für Produkte mit hohem Involvement.

Die Ergebnisse zeigen, dass das Involvement nicht als Moderator fungiert, das heißt der Grad des Involvements hat keinen signifikanten Einfluss auf die Konsumentenbewertung. Unabhängig vom Involvement erhöht die Offenlegung von KI jedoch die Authentizität und das Markenbild. Darüber hinaus wird eine signifikante Mediation von Authentizität auf das Markenbild und die Kaufabsicht festgestellt.

Diese Ergebnisse unterstreichen die Bedeutung wie KI in der Werbung kommuniziert wird, und deuten darauf hin, dass die Offenlegung die Wahrnehmung von Transparenz und Innovation fördern kann.

Schlagwörter: Authentizität, Elaboration Likelihood Model, Involvement, Kaufabsicht, KI, KI-generierte Werbung, Künstliche Intelligenz, Markenbild, Marketing, Werbung

Table of Contents

List of Abbreviations.....	V
List of Figures	VI
List of Tables.....	VI
1 Introduction	1
2 Theoretical Background	3
2.1 Artificial Intelligence.....	3
2.1.1 Definition of AI.....	3
2.1.2 Usage of AI in Marketing and Advertising.....	3
2.1.3 Challenges of AI.....	4
2.1.4 European AI Regulation.....	4
2.1.5 Consumers Evaluation of AI-Disclosure.....	5
2.2 Involvement.....	8
2.2.1 Elaboration Likelihood Model	8
2.2.2 Personal Involvement Inventory	10
2.3 Purchase Intention	11
2.3.1 Definition of Purchase Intention	11
2.3.2 Theory of Planned Behavior	12
2.4 Authenticity	12
2.5 Brand Image	13
3 Empirical Research	14
3.1 Research Design	14
3.2 Data Collection.....	16
3.2.1 Measures.....	16
3.2.2 Pre-Test	17
3.2.3 Sample.....	18
3.3 Data Analysis.....	18
3.3.1 Preparation & Descriptive Statistics	18
3.3.2 Assumption Testing.....	20

3.3.3	Testing Hypotheses	22
4	Discussion	27
4.1	Interpretation of Key Findings	27
4.2	Noteworthy Results	28
4.3	Strengths and Limitations	29
4.4	Future Research	30
5	Conclusion.....	32
	Bibliography.....	I
	Appendix	VI
	Appendix A: Survey	VI
	Appendix B: Pre-Test	IX
	Appendix C: Participants Demographics.....	X
	Appendix D: Main Study Results	XI
	Appendix E: Output for Hypothesis 1a	XVII
	Appendix F: Output for Hypothesis 1b	XIX
	Appendix G: Output for Hypothesis 1c	XXI
	Appendix H: Output for Hypothesis 2a	XXII
	Appendix I: Output for Hypothesis 2b	XXV

List of Abbreviations

AI	Artificial Intelligence
ANOVA	Analysis of Variance
CI	Confidence Interval
CTR	Click-Through-Rate
HC4	Heteroscedasticity-consistent Standard Error Estimate Type 4
PI	Purchase Intention
EU	European Union
ELM	Elaboration Likelihood Model
MWh	Megawatt-hour
ns	Non-significant
PII	Personal Involvement Inventory
PBC	Perceived Behavioral Control
tCO ₂ e	Tonnes of Carbon Dioxide Equivalent
VIF	Variance Inflation Factor
Wh	Watt-hour

List of Figures

Figure 1: Elaboration Likelihood Model	9
Figure 2: Conceptual Framework	15
Figure 3: Graph Authenticity/Disclosure	23
Figure 4: Graph Brand Image/Disclosure.....	24
Figure 5: Results of H1	24
Figure 6: Results of H2.....	26

List of Tables

Table 1: Study results for consumer evaluations of AI-disclosure.....	7
Table 2: Low and High Involvement Criteria	10
Table 3: Descriptive statistics of the variables	19
Table 4: ANOVA: Authenticity	22
Table 5: Moderation Analysis Disclosure x Involvement for Brand Image.....	23
Table 6: Test(s) of highest order unconditional interaction(s)	23
Table 7: Indirect effect(s) of X on Y	25
Table 8: Summarized Mediation Table for Brand Image.....	25
Table 9: Summarized Mediation Table for Purchase Intention.....	26

1 Introduction

‘Please write a fairytale for a five-year-old girl, including unicorns and rainbows.’ - whether producing creative content, ask simple search queries or generating complex programming codes, the possibilities of Artificial Intelligence (AI) seem endless. Hence, nowadays AI is embedded in everyday life and tools are also increasingly adopted in marketing (Lim & Schmälze, 2024, p. 1). For instance, DALL-E from OpenAI, generates visual assets from text prompts, while other tools assist with copywriting, campaign planning, and targeting strategies (Feuerriegel et al., 2024, p. 114; Wu et al., 2025, p. 21). Creating advertisements involves high costs and time investments, including scripting, photo shooting, editing, and media placing (Feuerriegel et al., 2024, p. 120). AI offers the opportunity to reduce these costs and accelerate the creative process (Grigsby et al., 2025, p. 1). Companies such as KitKat has already publicly disclosed their use of AI (Kučinskas & Survilaite, 2025, p. 727). However, it is likely that many other brands employ AI as well but without labelling it explicitly - particularly given the difficulty in distinguishing AI-generated from human-created content (Karpinska-Krakowiak & Eisend, 2024, p. 6; Lim & Schmälze, 2024, p. 1).

Currently, such disclosure is voluntary. This will change with the upcoming AI Act of the European Union (EU), which mandates that all AI-generated content be clearly labelled from August 2026 onward (Regulation (EU) 2024/1689, 2024, pp. 44ff). When companies do this for advertising purposes, it raises several questions regarding consumer evaluations - the thesis will examine the following:

RQ1: How do consumers respond to advertisements when they are informed that the content is AI-generated - does transparency affect the perceived authenticity, brand image and purchase intention?

Prior research offers conflicting insights. Several studies indicated that disclosure led to more negative perceptions, particularly when AI was explicitly disclosed (Grigsby et al., 2025, p. 1; Karpinska-Krakowiak & Eisend, 2024, p. 7; Lee & Kim, 2024, p. 8). For example, Kučinskas & Survilaite reported a reduction in purchase intention and brand trust (2025, p. 727) and Baek et al. discovered a decrease in ad attitude and ad credibility (2024, p. 8). In contrast, Kirkby et al. found no significant effect of disclosure on either authenticity or brand attitude (2023, p. 1115). These mixed findings highlight the need for further research into the role of AI disclosure.

Additionally, the thesis considers whether the impact of AI disclosure varies depending on the type of product being advertised. This distinction is grounded in the Elaboration Likelihood

Model (ELM) by Petty and Cacioppo (1981), which posits that consumers process information differently depending on their level of involvement. Low involvement products are typically of limited personal relevance (Wong & Sheth, 1985, p. 381), whereas high involvement products are evaluated more critically (Petty & Cacioppo, 1986, p. 125). This leads to the second research question:

RQ2: Do the effects of AI disclosure on perceived authenticity, brand image, and purchase intention differ between advertisements for low and high involvement products?

This specifically addresses the current research gap, as previous studies did not differentiate by the level of involvement, they rather focused on present or absent AI disclosure. As such, this research offers much-needed insights into how consumer perceive AI-generated advertisement. To investigate these questions, this study employs a 2x2 between-subjects experimental design, implemented via an online survey, manipulating both AI disclosure and product involvement. The findings aim to provide both theoretical and practical insights for marketers preparing for a future legal requirement to disclose AI usage. The study helps clarify whether disclosure harms or enhances consumer perception and offers insights into how AI should be communicated in brand strategy.

2 Theoretical Background

The theoretical background provides the conceptual foundation of this study, beginning with an overview of Artificial Intelligence, followed by the concept of involvement, and continuing with an overview of the dependent variables purchase intention, authenticity, and brand image.

2.1 Artificial Intelligence

The subsection on Artificial Intelligence gives an overview of the technology, its usage in marketing, consumers evaluations of its application, and current challenges of AI including the European regulatory framework.

2.1.1 Definition of AI

AI is a machine-based system that can generate output from the input it receives (Regulation EU, 2024, Article 3 (1)). This input is being used for machine learning, which is a subset of AI. Instead of fully programming the system, the algorithm independently analyzes and interprets the data to improve its accuracy (Haleem et al., 2022, p. 119). These algorithms, also called generative AI (GenAI) can then generate new data such as content, predictions or recommendations (Feuerriegel et al., 2024, p. 111f). Examples for these kinds of technologies are DALL-E 2, Copilot or GPT-4. Also, companies like Google or Meta provide such systems (Grigsby et al., 2025, p. 1; Lim & Schmälze, 2024, p. 1).

2.1.2 Usage of AI in Marketing and Advertising

Feuerriegel et al. perceive this technological advancement as a “revolutionary opportunity” (2024, p. 111). 91.8 % of 170 marketing managers surveyed agree, as they consider the implementation of AI as important (Bünthe, 2023, p. 10). They use it most for the marketing execution and second most for performance marketing (Bünthe, 2023, p. 14). Examples of use are interactive website design, content creation, intelligent emails, and many more possible applications in digital marketing (Haleem et al., 2022, p. 119). AI can also create ‘special’ content - the so called ‘synthetic advertisement’. This is a form of manipulated ad, whereas data is being modified and personalized like sunglasses ad with the consumers’ own face (Campbell, 2022, p. 22). In general, GenAI makes the marketing environment more efficient and less costly (Grigsby et al., 2025, p. 1).

Various popular brands like Heinz openly communicate their utilization of AI to earn consumer trust and to differentiate the brand (Kučinskas & Survilaite, 2025, p. 727).

2.1.3 Challenges of AI

Artificial Intelligence has many advantages like the easy generation of output; however, this output is prone to error or the production of high carbon emissions in the training process. A study by Google and the University of California at Berkeley calculated the carbon emissions of the training of GPT-3, as this process consumes most energy (de Vries, 2023, p. 2191). They estimated 552tCO₂e, which is three times the emissions of an entire passenger jet flying between San Francisco and New York. The energy consumption of training the AI takes 1,287 MWh, whereas a Bitcoin purchase costs about 0.7 MWh (Patterson et al., 2021, p. 13). In terms of energy consumption per request, a ChatGPT enquiry uses almost 10 times more energy than a standard Google search, equating 2.9 Wh (de Vries, 2023, p. 2193). However, AI may also produce incorrect outputs or violates copyright laws, especially when the quantity and quality of training data is low (Feuerriegel et al., 2024, p. 117). This raises ethical questions, as it not yet clear who is responsible for wrong or illegal results, the AI or the human. In addition, privacy protection must also be ensured for ethical correctness because data collection becomes an integral part of the technology (Campbell et al., 2022, p. 33). Beyond these concerns, AI poses challenges related to the job market. It has the potential to replace various jobs, including roles in creative fields such as writing, design, and marketing. For example, Duolingo, a language learning app, just announced to be “AI-first”, meaning that they “stop using contractors to do work that AI can handle” (Duolingo, 2025). This raises questions about which advertising tasks will still require human creativity, supervision, and intervention (Huh et al., 2023, p. 480).

2.1.4 European AI Regulation

In June 2024, the European Union introduced the world's first comprehensive AI law, the so-called ‘Artificial Intelligence Act (EU Regulation 2024/1689)’, short ‘AI Act’, to ensure a safe and transparent application (European Parliament, 2025, p. 5; Regulation (EU) 2024/1689, 2024, p. 1). It categorizes AI systems depending on their risk levels and sets rules for the use and disclosure of these systems across member states. However, GenAI is not considered as high-risk (European Parliament, 2025, p. 3), but needs to fulfil transparency requirements defined in Article 50 of the AI Act. Paragraph 4 defines AI-disclosure as following: “Deployers of an AI system that generates or manipulates image, audio or video content constituting a deep fake, shall disclose that the content has been artificially generated or manipulated” (Regulation (EU) 2024/1689, 2024, Article 50 par. 4).

Additionally, regulations like the prevention of creating illegal content as well as the publication of the training data used must be implemented by August 2026 at the latest (Regulation (EU) 2024/1689, 2024, p. 44).

2.1.5 Consumers Evaluation of AI-Disclosure

Different studies have examined the influence of disclosure of AI in advertisements on consumer behavior although results are affected by the AI tool used, sample demographic, and technological advancement of the participants (Mehta et al., 2022, p. 2014).

Grigsby et al. researched the perceived AI-labelling of service ads with differentiation between intangible and tangible attributes. Results show that, trust and attitudes toward the ad are more positive when the ad is not labelled as AI-generated. Instead, AI-disclosure results in lower trust as well as ad attitude, especially for intangible attributes. Therefore, these should not be used in service advertisements, whereas tangible attributes can be created by AI without affecting consumers' perceptions (Grigsby et al., 2025, p. 1-4).

Similar results found Kučinskas & Survilaitė as they suggest integrating human elements to the AI-generated content. This is because AI disclosure negatively affects purchase intention, brand image, and trust, which is mediated by perceived authenticity. Consumer attitudes even worsen if established brands, compared to unknown one's, disclose their AI usage (Kučinskas & Survilaitė, 2025, p. 727). Additionally, Lee & Kim found that authenticity mediates the expected product quality (2024, p. 8). Thereby, marketers need to manage consumers perception of authenticity when using AI in advertising (Kučinskas & Survilaitė, 2025, p. 732). Consumers show a higher willingness to discuss an advertisement when it is disclosed that the ad was placed by AI, compared to when it is disclosed that the ad was created by AI. This tendency becomes stronger when consumers believe AI is capable of performing complex tasks (Wu et al., 2025, p. 31).

These findings indicate that the impact of AI disclosure depends on the extent of how much AI is involved. In relation, Karpinska-Krakiwiak and Eisend examined consumers perception with different levels of disclosure, whereas they discovered that the type of disclosure affects consumers attitudes. Specifically, simple disclosures lead to more positive attitudes compared to extended disclosures (2024, p. 7). Supporting this, Baek et al. found that AI disclosure significantly reduced both ad credibility and ad attitude when compared to the same ad without disclosure (2024, p. 8).

Interestingly, Wu et al. did not find a significant effect of disclosure on word-of-mouth intention, without controlling for complexity perceptions of AI (2025, p. 27). Similarly, Kirkby et al. found no significant (ns) relationship between brand voice authenticity and disclosure

condition ($b = -.071, p = .161$), nor between disclosure and brand attitude ($b = .028, p = .382$) (2023, p. 1115). However, their study revealed that brand authenticity positively affects brand attitude, suggesting a potential mediating role of authenticity shaping consumer evaluations ($b = .837, p = .0000$) (Kirbky et al, 2023, p. 1115).

Many other studies also examined consumer perceptions of AI disclosure in contexts other than advertising, using scenarios such as hotel reviews or sales calls instead. Hotel reviews were not disclosed as AI but if it seems to be created by AI people rate the usefulness, authenticity and trustworthiness much lower (Amos & Zhang, 2024, p. 7). Arango et al. show similar results for the donation intention when images of non-existent children are generated by AI and labelled as such (2023, p. 493). It can be assumed that this will also apply to the purchase intention of falsified advertising images. Lee & Kim found only marginally significant evidence for a decrease in purchase intention (PI) when AI is involved ($p = .058$), but they report a strongly significant reduction in perceived authenticity ($p < .001$) (2024, p. 5). In contrast, Luo et al. found a significant decline in PI when customers are aware that they are interacting with AI ($p < .01$), with PI dropping by nearly 80 % as well as shortening the duration of the call (2019, p. 942). However, in terms of purchase intention, AI-chatbots not labelled as AI are as effective as skilled agents and four times more successful than those without experience (Luo et al., 2019, p. 937). AI not only increases effectiveness but also Click-Through-Rates (CTR) as the study of Xia et al. showed: advertisements in a field experiment (Instagram and Facebook) labelled as AI-generated held higher CTR than the same ad labelled as human-generated (2025, p. 10). This shows that AI-disclosure increases clicks in social media. However, people rate AI-generated messages as more effective than human-generated ones, but only when the source is not revealed. Once it is disclosed, effectiveness is rated the same (Lim & Schmälzle, 2024, p. 8), whereas not only this study shows a general bias against AI-generated content. For instance, Lee & Kim also found negative authenticity responses towards AI versus human designed clothes (2024, p. 8). An explanation for negative disclosure effects might be the discomfort of humans towards machines (Luo et al., 2019, p. 937). Therefore, marketers need to “carefully weigh the potential benefits”, which are a higher efficiency and lower cost (Grigsby et al., 2025, p. 1), “against the possible negative impact of disclosing AI” (Wu et al., 2025, p. 33). However, disclosure could also have a positive effect on consumers by enabling them to make informed choices (Karpinska-Krakowiak & Eisend, 2024, p. 9). Ultimately, when deciding to use AI, advertisers will need to follow a strategy of either delaying disclosure, addressing only AI experienced customers (Luo et al., 2019, p. 945) or openly disclosing their use of AI to demonstrate transparency and innovation (Kučinskas & Survilaite, 2025, p. 727).

Table 1: Study results for consumer evaluations of AI-disclosure

Author(s)	Year	AI Disclosure	Context of Study	Variable(s)	Effect
Amos & Zhang	2024	Non-disclosed, but perceived AI-generated	Hotel reviews perceived as AI-generated	Usefulness Authenticity Trustworthiness	↓ ↓ ↓
Arango et al.	2023	Disclosed	Donation intention when AI-generated images of nonexistent children	Donation intention	↓
Baek et al.	2024	Disclosed	AI-generated online advertising message with present vs. absent disclosure label	Ad credibility Ad attitude	↓ ↓
Grigsby et al.	2025	Disclosed	Service advertisements to & use of intangible vs. tangible attributes	Trust Ad attitude	↓ ↓
Kirkby et al.	2023	Disclosed	Marketing text (product description & chatbot conversation) disclosed as AI, 'human' or not disclosed	Authenticity Brand attitude	ns ns
Kučinskas & Survilaite	2025	Disclosed	Ads with a well-known brand vs. fictitious brand	Authenticity Purchase intention Brand trust Brand image	↓ ↓ ↓ ↓
Lee & Kim	2024	Disclosed	AI-generated vs. human-designed fashion in the fashion industry	Purchase intention Authenticity Product attitude	↓ ↓ ↓
Lim & Schmälzle	2024	Non-disclosed	Message (for vaping prevention) AI vs. human generated	Perceived effectiveness	↑
		Disclosed		Perceived effectiveness	—
Luo et al.	2019	Disclosure	Customer Service Sales Calls (chatbot vs. human agents)	Purchase intention	↓
Wu et al.	2025	Disclosed	Ad created by human but disclosed as AI-generated & ad was chosen by AI	WOM intention Ad attitude Brand attitude Purchase Intention	— — ↓ —
Xia et al.	2025	Disclosed	AI-generated product packaging vs. human-generated	Brand modernity Brand evaluation CTR	↑ ↑ ↑

2.2 Involvement

The section on involvement introduces the concept through the Elaboration Likelihood Model and makes it measureable using the Personal Involvement Inventory.

2.2.1 Elaboration Likelihood Model

The Elaboration Likelihood Model (ELM) developed by Petty and Cacioppo in 1981, describes how people process incoming information (Petty & Cacioppo, 1986, p. 125). This framework is also used in marketing to explain how advertising is processed, in which people become differently involved and in which consumers attach meaning to a product category (Baati & Akrouf, 2024, p. 736). The individual involvement of a person depends on the type of media used, their personal situation like the individual evaluation of the product as well as their characteristics such as the prior experience or knowledge (Petty & Cacioppo, 1984, p. 674; Zaichkowsky, 1986, p. 5). These factors lead to either low involvement of a person, which is processed via the peripheral route or high elaboration likelihood via the central route (Petty & Cacioppo, 1984, p. 673). Figure 1 provides a schematic overview of how messages are processed via the central or peripheral route. In the peripheral route, the motivation and/or ability to process persuasive messages is relatively low because it is mostly associated with positive or negative cues. Non-message factors like the expertise or attractiveness of the source are more important than the presentation of relevant information (Petty & Cacioppo, 1984, p. 673). In contrast, highly involved people carefully consider the quality of the message and the information presented because it is personally relevant to them, but they also actively seek product- and market information (Petty et al., 1983, p. 136; Zaichkowsky, 1986, p. 6). Furthermore, the ability to process an information depends on factors like the presence of distractions, repeated views, and prior knowledge. As ability increases, so does cognitive activity, leading to an intensive processing. Those cognitive responses can be either positive, negative or neutral (Petty & Cacioppo, 1986, p. 4).

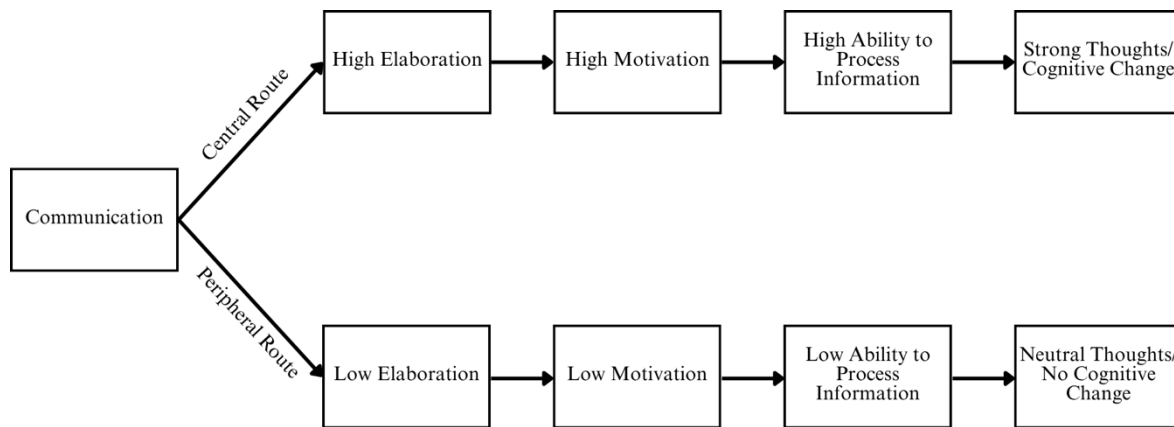


Figure 1: Elaboration Likelihood Model

Note. Adapted from *Communication and persuasion. Central and peripheral routes to attitude change* (p. 4), by R. E. Petty & J. T. Cacioppo, 1986, Springer Science and Business Media.

Price and risk are typical factors, that categorize the involvement level. Low involvement products are usually cheap and yield little risk (Jain, 2019, p. 944), while high involvement products are expensive, therefore also higher in financial risk but also bear social and psychological risk (Wong & Sheth, 1985, p. 381). Hence, people collect much information from various sources to minimize the risks, whereas low involvement are routine purchases (Jain, 2019, p. 945). Consumers are willing to put much energy and time into the careful consideration of this information, which makes the high involvement process itself complex, but also complex in terms of persuasion (Petty & Cacioppo, 1986, p. 125). In contrast, limited effort put in the decision-making process account for low involvement (Wong & Sheth, 1985, p. 381). This shows the importance of the level of information needed for the consumer's purchase decision (Jain, 2019, p. 943).

Different marketing strategies such as specific communication approaches or the provision of detailed information are required depending on the level of involvement (Jain, 2019, p. 943). Although there are typical product categories that represent low or high involvement goods, it is still a person's individual determination of the relevance of a particular product. One person may be highly involved with a sports product because he or she is a supporter of a sports team, where another person is not. Thus, the same product represents another involvement level for different people (Kirkby et al., 2023, p. 1116), which makes it difficult for marketers to address. Additionally, the level of involvement depends on the identified need consumers have and how motivated they are to solve the problem (Jain, 2019, p. 947).

Table 2 compares the typical attributes for both level of involvements.

Table 2: Low and High Involvement Criteria

Low Involvement (Peripheral Route)	High Involvement (Central Route)
Low personal relevance & low motivation	High personal relevance & high motivation
No extensive thoughts about the product	Many information required & evaluated
Usually cheap & bought regularly	Usually expensive & bought rarely
Decision bears little risk	Decision includes higher financial, social or psychological risk
Example: Sunscreen (Akbari, 2015, p. 483)	Example: Laptop (Akbari, 2015, p. 473)

Many studies use the level of involvement as a moderator, since consumer perceptions may differ between involvement levels. For instance, Eisend investigated the role of involvement as a moderator for the persuasiveness of message sidedness and found that the negativity of information was more persuasive under high involved consumers, while the amount of information was more persuasive under low involvement (2013, p. 567). Barari et al. highlight the moderating role of involvement by showing that the ‘dark side’ of AI, like privacy concerns or perceived risks, can negatively influence customer responses such as purchase behavior, loyalty, and well-being. Their results indicate that these adverse effects are more pronounced for high-involvement products compared to low-involvement ones (2024, p. 1247).

Based on these findings, the present research uses involvement as a moderator as well, which leads to the following hypothesis:

Hypothesis 1: Disclosure of AI in ad creation decreases perceived authenticity, brand image and purchase intention. This effect is moderated by the level of involvement, such that ads of high involvement products are perceived worse.

2.2.2 Personal Involvement Inventory

“Different people are differently involved with the same product” (Zaichkowsky, 1994, p. 62), therefore Zaichkowsky developed a construct known as the Personal Involvement Inventory (PII) to measure individual involvement (1985, p. 341). Initially, the PII consisted of 20 items presented on a 7-point semantic differential scale, using adjective pairs such as ‘*irrelevant/ relevant*’ and ‘*uninterested/ interested*’ where the left end indicated low involvement and the right end indicated high involvement. Totaling the items gives a score from 20 to 140, whereas a product is defined as ‘low involvement’ with a score from 20 to 69 and high involvement from 111 to 140 (Zaichkowsky, 1985, p. 350). Studies showed that products such as instant coffee (66) and bubble bath (69) account for low involvement, while an automobile (122) is associated with high involvement (1985, p. 350).

Ten years later, Zaichkowsky revised the scale, reducing it to 10 items while maintaining high reliability - originally reporting a Cronbach's Alpha of 0.95 and still achieving a reliability of 0.90 (1994, p. 60f). Notably, all but one item from the revised scale were also present in the original version (Zaichkowsky, 1994, p. 67). The updated scale is divided into two subscales: five items measuring cognitive involvement and five items measuring affective involvement (Zaichkowsky, 1994, p. 66). The items are (Zaichkowsky, 1994, p. 70):

- I. important/unimportant (reverse coded)
- II. boring/ interesting
- III. relevant/ irrelevant (reverse coded)
- IV. exciting/ unexciting (reverse coded)
- V. means nothing/ means a lot to me
- VI. appealing/ unappealing (reverse coded)
- VII. fascinating/ mundane (reverse coded)
- VIII. worthless/ valuable
- IX. involving/ uninvolved (reverse coded)
- X. not needed/ needed

This scale was later adapted by Celuch & Taylor, who developed an 8-item version removing the items '*interesting*' and '*involving*' (1999, p. 115).

2.3 Purchase Intention

One of the dependent variables in this study is purchase intention, wherefore it will be explained in the following section.

2.3.1 Definition of Purchase Intention

Generally, intention is "the degree to which a person resolves to act in a certain way" (Morwitz & Munz, 2021, p. 26) and Spears & Singh define the purchase intention (PI) as "an individual's conscious plan to make an effort to purchase a brand" (2004, p. 56).

The simplest approach to measure the PI is to ask a person about their intention to buy a product (Ajzen & Fishbein, 1980, p. 5; Wong & Sheth, 1985, p. 378). Nevertheless, these intentions are not a thoroughly reliable indication for subsequent behavior as many scholars indicate (Morwitz & Munz, 2021, p. 31; Wong & Sheth, 1985, p. 378). Marketers call this phenomenon the intention-behavior-gap (Conner & Norman, 2022, p. 2). Although this difference exists, the PI is an important key figure and "a good proxy for the ultimate behavior [organizations] are interested in" (Morwitz & Munz, 2021, p. 26)

2.3.2 Theory of Planned Behavior

The theory of reason action developed by Fishbein and Ajzen in 1967 is a framework to understand why people behave differently. This addresses the intention-behavior-gap and the variance of intention and actual behavior (Ajzen, 1991, p. 171).

According to the theory, a person's intention can be divided into two parts:

1. Attitude toward the behavior
2. Subjective norm

The attitude towards the behavior reflects a person's own feeling whether this behavior is good or bad (Ajzen & Fishbein, 1980, p. 6). People having a favorable attitude towards a product or brand, the purchase intention will be higher (Ajzen, 1991, p. 189).

But not only the personal judgment plays a role in the buying decision, but also the perceived societal judgment towards this behavior - the subjective norm (Ajzen & Fishbein, 1980, p. 6). It mirrors the "perceived social pressure to perform or not to perform a specific act" (Ajzen, 1991, p. 188). Purchase intentions usually align with subjective norms to be accepted by society (Ajzen, 1991, p. 188).

The theory of planned behavior extends this theory with a third component that influences a person's behavior, which is the perceived behavioral control (PBC). The PBC describes a person's perception of the difficulty to perform a certain behavior and their ability to do so. A higher PBC leads to higher PI (Ajzen, 1991, p. 181).

2.4 Authenticity

In recent years, the term authenticity has become increasingly important, particularly in management journals, where the number of related articles has more than doubled (Lehman et al., 2019, p. 1). Although the term is widely used, its definition varies across the literature, however scholars constantly refer to "real", "genuine" and "true" (Beverland & Farrelly, 2010, p. 838). Whereas, in everyday discourse authenticity is the opposite of 'falsity' or 'fakery' (Dutton, 2003, p. 258f). As Dutton (2003, p. 258) observes, authenticity functions as a 'dimension word' - its meaning is contingent upon the specific aspect of the referent being discussed. So, the critical question is: authentic in what respect, or true to what? Therefore, each situation "involves the application of a different meaning of the concept [authenticity]" (Lehman et al., 2019, p. 2). Lehman et al. reviewed 327 management journal articles to better define authenticity. In the end, they found three different perspectives on authenticity and subdivided them into three categories: consistency, conformity, and connection (2019, p. 4f). According to consistency, "an entity is authentic to the extent that it is consistent in terms of its external expressions on the one hand, and its internal values and beliefs on the other hand"

(Lehman et al., 2019, p. 5). If the entity conforms to the social category to which it claims to belong, it is authentic in terms of conformity (Lehman et al., 2019, p. 12). The third perspective of authenticity, connection, is present when an entity is connected to a person, place, or time as claimed (Lehman et al., 2019, p. 16). In the context of marketing, it has been found that consumers all seek authenticity in products and brands, but for different reasons such as control, connection and/or virtue (Beverland & Farrelly, 2010, p. 853). A brand is perceived as authentic when the communication and behavior is transparent, consistent, and genuine. This leads to effective ads and positive consumer responses like brand attitude and purchase intention (Kučinskas & Survilaite, 2025, p. 729; Lee & Kim, 2024, p. 3). Consequently, authenticity acts as a mediator in most consumer behavior scholars (Kučinskas & Survilaite, 2025, p. 727; Lee & Kim, 2024, p. 5) - this is also hypothesized for the present study:

Hypothesis 2: Authenticity mediates the effect of AI disclosure on purchase intention and brand image.

2.5 Brand Image

The term brand image is relevant in the consumer behavior research but there is no clear definition as the phrase has grown and evolved in the literature over time. Dobni & Zinkham tried to group different definitions into five categories: the “*blanket definitions*”, which are rather abstract and not very well defined, the “*emphasis on symbolism*”, so the consumer has a personal or social connection to the product, while the “*emphasis on meanings or messages*” focuses on the meaning a consumer connotes with the product. Associating the self-perception with the image of a brand is categorized as the “*emphasis on personification*”, and finally the “*emphasis on cognitive or psychological elements*” focuses on the consumer’s process when getting in touch with the brand (1990, p. 112ff). Ultimately, Dobni & Zinkham define the brand image broadly as “the concept of a brand that is held by the consumer” (1990, p. 118), regardless of their rational or emotional interpretation. In general, it is a subjective perception of the feelings and associations towards the brand, usually shared by a group of consumers (Riezebos & Riezebos, 2003, p. 63).

3 Empirical Research

The third chapter presents the empirical research of this study, beginning with the research design, followed by the data collection process, and concluding with the data analysis.

3.1 Research Design

This study adopted a quantitative research approach, using a laboratory survey to investigate the impact of AI disclosure in advertising across different levels of product involvement. The experiment followed a 2x2 between-subjects design, so participants were randomly assigned to one of the four conditions (disclosure + low involvement/ disclosure + high involvement/ non-disclosure + low involvement/ non-disclosure + high involvement). Thus, each participant saw one advertisement and answered questions regarding the perceived authenticity, brand image and purchase intention afterwards. The only difference was that participants from the experimental group received explicit information that the advertisements were generated by AI, whereas those in the control group did not receive such disclosure. Additionally, the presented advertisements differed depending on low or high involvement condition. Data was collected in the form of an online experiment, using the software SoSci Survey.

The present study used hypothetical brands as an advertising stimulus, because people may already hold positive or negative associations with a brand. Using non existing brands increases the internal validity (Geuens & De Pelsmacker, 2017, p. 86). A preliminary study was conducted to identify suitable low- and high-involvement products for the advertisements, ensuring that the selected items aligned with the involvement levels of the sample group, which increased internal validity and ensured that the results reliably represent the involvement levels. Extracts from Zaichkowsky's Personal Involvement Inventory (PII) were utilized to assess the individual involvements using four products - two representing low and two high involvement. The PII consists of 20 semantic differential items (e.g., "*matters to me/does not matter to me*", "*unexciting/exciting*") (Zaichkowsky, 1985, p. 350), whereas only four were used.

After choosing the products, ads were generated with ChatGPT and Canva. The ads were rather simple and realistic instead of creative and sophisticated to prevent confusion or distraction from the stimulus (Geuens & De Pelsmacker, 2017, p. 85). Realism of the ad also contributed to the fact that participants of the control group did not recognize the AI generation, which might distort the results. The advertisements shown to the participants can be found in appendix A. Generally, the ads were similar in terms of font, position and color, only the fictional brands and pictured products differed. To increase external validity, the disclosure label mirrored those used by platforms such as Meta and TikTok. These platforms use "Made with AI" or "AI Info", but participants saw the following notice at the top of the ad: *'This ad was created with Artificial*

Intelligence.’ (Wu et al., 2025, p. 26). This longer label and positioning ensured that participants clearly recognized the use of AI. By allocating the participants randomly to a group the external validity increased as well (Kuß et al., 2014, p. 187). At the end of the survey, an attention check was added, asking the participants what product the advertisement showed, to verify their engagement. If the answer was incorrect, the participants were removed to guarantee the quality of the data analyzed. A manipulation check was also included, asking participants if they perceived the ad as AI-generated. Additionally, demographic measures such as gender, sex, and education were collected to have greater knowledge about the sample group and to draw possible conclusions depending on these characteristics. The survey is attached in appendix A. This methodology provides important insights into how consumers perceived AI-generated and -labeled advertisements, particularly in terms of perceived authenticity, brand image, and purchase intention, while also differentiating between low and high involvement products. Thereby, the level of involvement functioned as a moderator (Hayes, 2018, p. 221).

As mentioned in the theory part, scholars found that authenticity positively affected consumer responses, such as purchase intention, brand trust or brand/ad attitude (Baek et al., 2024, p. 8; Lee & Kim, 2024, p. 5). Kučinskas & Survilaitė even discovered the mediating negative effect of authenticity on consumers evaluations when AI disclosure was present (2025, p. 727). Hence, this study assumes that authenticity acts as a mediator.

Figure 2 illustrates the conceptual framework. AI Disclosure and Involvement are independent variables; authenticity, brand image and purchase intention the dependent ones.

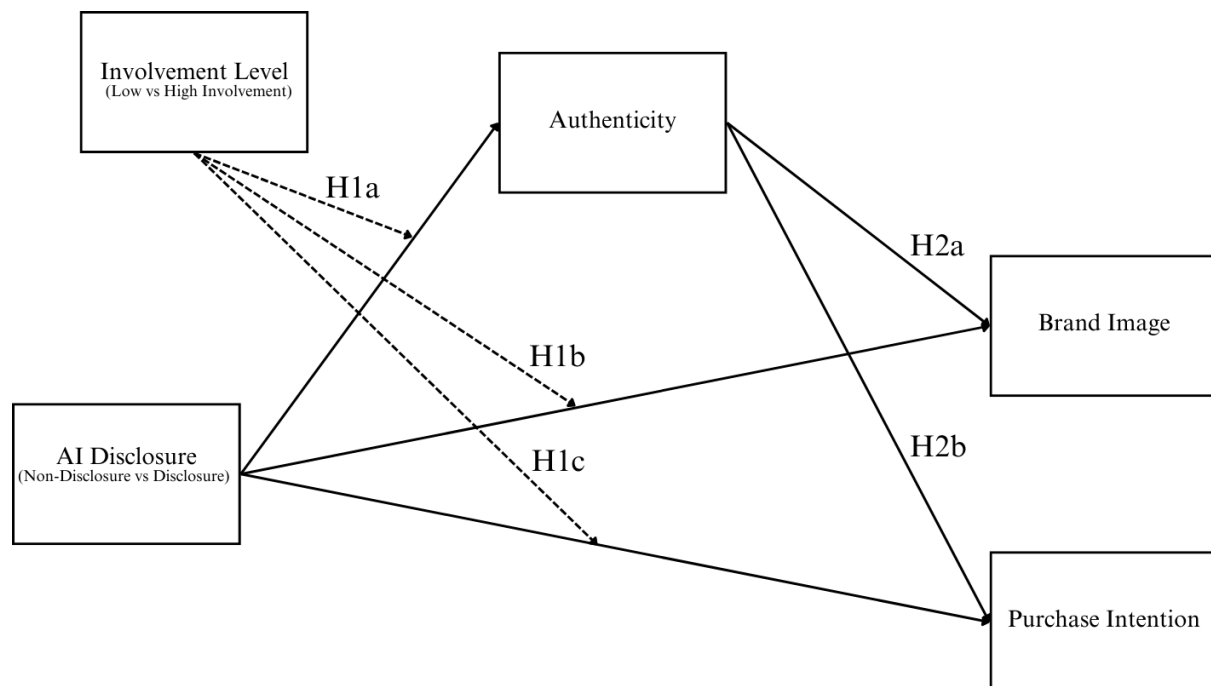


Figure 2: Conceptual Framework

The following research hypotheses will be examined.

Hypothesis 1: Disclosure of AI in ad creation decreases perceived authenticity, brand image and purchase intention. This effect is moderated by the level of involvement, such that ads of high involvement products are perceived worse.

H1a) Disclosure of AI usage reduces perceived authenticity. This effect is stronger for ads of high than low involvement products.

H1b) Disclosure of AI usage reduces the brand image. This effect is stronger for ads of high than low involvement products.

H1c) Disclosure of AI usage reduces purchase intention. This effect is stronger for ads of high than low involvement products.

Hypothesis 2: Authenticity mediates the effect of AI disclosure on purchase intention and brand image.

H2a) The effect of disclosure on brand image is mediated by perceived authenticity.

H2b) The effect of disclosure on purchase intention is mediated by perceived authenticity.

3.2 Data Collection

The data collection section begins the measurements for the variables, followed by the pre-test procedure, and concluding with the sample distribution of the main study.

3.2.1 Measures

The study examined variables such as authenticity, brand image, and purchase intention. All items were measured using a 5-point semantic differential scale. To evaluate brand image and PI, items from Spears and Singh (2004, p. 60) were employed, (2004, p. 60). PI was measured with a five item scale and the question “How likely are you to purchase the product after seeing the ad?” by providing answers such as “Probably not buy/ probably buy”. “I find the brand...” was the statement used to determine the brand image. The five items assessed are bipolar, consisting of *unappealing/appealing* (1), *bad/good* (2), *unpleasant/pleasant* (3), *unfavorable/favorable* (4) and *unlikeable/likeable* (5).

To examine the authenticity a five item semantic differential scale derived from Beverland and Farrelly (2010) was used. Respondents were asked to express the perceived authenticity based on the advertisement by rating five items, such as *ingenuine/genuine* (1), *not authentic/authentic* (2), *not original/original* (3), *unfaithful/faithful* (4) and *unnatural/natural* (5).

Parts of the reduced ten item scale proposed by Zaichkowsky (1994, p. 65) were used to measure involvement. Celuch & Taylor reduced them to the eight most reliable and valid items, so they removed “interesting” and “involving” (1999, p. 115). The scale used for this survey was further reduced to two affective and two cognitive dimensions (Celuch & Taylor, 1999, p. 117), sticking for this research with *unexciting/exciting* (1), *appealing/unappealing* (2), *means a lot to me/means nothing to me* (3) and *valuable/worthless* (4). All items can be found in the appendix.

3.2.2 Pre-Test

A pre-test with N = 11 was conducted to test the survey on its understanding as well as to select the most relevant stimuli for the main study. As described above, Zaichkowsky’s PII was used to assess the degree of involvement participants associate with various products. According to the scale, it is a low involvement product between 20 and 69, and high involvement from 111 to 140. Bubble bath (score of 69) and automobile (score of 122) were already indicated as low/high involvement (Zaichkowsky, 1985, p. 351), which was used in the pre-test as well. Additionally, toothpaste and laptop were used to assess the involvement degree as this has been used in the literature as stimuli as well (Akbari, 2015, p. 478; Petty et al., 1983, p. 139).

Low Involvement

- Score Product 1 (Bubble Bath): 17
- **Score Product 2 (Toothpaste): 15**

High Involvement

- Score Product 3 (Automobile): 22
- **Score Product 4 (Laptop): 22**

A paired samples t-test revealed a significant difference in the involvement of toothpaste and laptop ($p = .026$), therefore these products were appropriate for this research. A laptop was chosen instead of a car because the study primarily involved young participants, who are likely to feel a stronger connection to laptops than cars. Therefore, the study used toothpaste as the low involvement product and laptop as the high involvement product. Next to the involvement assessment, participants tested the questions of the survey on its understanding by providing an open text input field.

The Cronbach’s Alpha was calculated based on the results of the pre-test, yielding values of .836 for authenticity, .936 for brand image, and .95 for purchase intention. These indicated strong internal consistency for all scales. However, participants reported that the items measuring purchase intention were overly similar. Therefore, this scale was replaced with a single-item measure. No other adaptations needed to be made.

3.2.3 Sample

Convenience sampling was used as the sampling procedure by employing participants via personal contacts, and social media. The data was collected in a two-week time span, ranging from the 8th of May to the 19th of May 2025. A total of 137 people finished the survey. After removing those who did not pass the attention check, a total of 136 participants remained. The study aimed for approximately 30 participants per group, based on this number being sufficient to detect potential differences between conditions (Geuens & De Pelsmacker, 2017, p. 87). The final sample for this study consisted of participants randomly allocated to experimental conditions, with group sizes ranging from 29 to 36 individuals. 53.5 % out of the 72 participants in the control group perceived the ad as AI-generated, which might have biased their answers. The majority of participants were female (68.4 %), followed by male participants (30.9 %), and other (0.7 %). In terms of generational distribution, the majority belonged to Generation Z (14 to 29 years old), accounting for 72.1 % of the total, followed by Generation X (45 to 68) at 17.6 %, and Generation Y (30 to 44) with 10.3 %.

Regarding educational attainment, over 40 % held a bachelor's degree, one-fourth completed a master's degree, and approximately 20 % obtain A-level qualifications. Therefore, the study reflected primarily young, well-educated individuals. A detailed overview of the demographics is included in appendix C. The data were collected anonymously, however informed consent was obtained to ensure ethical compliance.

3.3 Data Analysis

The following section presents the data analysis, beginning with the preparation of the dataset and descriptive statistics, followed by assumption testing, and concluding by testing the hypotheses.

3.3.1 Preparation & Descriptive Statistics

The dataset was extracted from SoSci Survey and prepared for further analysis. Reverse-coded items within the Involvement scale were recoded to ensure consistent interpretation. New variables were created by calculating mean values for the constructs authenticity, brand image, and involvement. A binary variable *Disclosure* was created to represent the experimental condition: participants were coded as "disclosed" (1) if the stimulus variable *IMOI* equaled 1 or 3, and as "not disclosed" (0) otherwise. Additionally, a binary variable for Involvement level was generated to differentiate between product types: participants who saw the low-

involvement product (toothpaste) were coded as 0, and those who saw the high-involvement product (laptop) as 1. This processed dataset served as the basis for the following analyses.

Participants in the disclosure condition reported slightly higher perceived authenticity ($M = 2.59$, $SD = 0.89$) compared to those in the non-disclosure condition ($M = 2.33$, $SD = 0.90$). A similar pattern was observed for brand image, where evaluations were more favorable in the disclosure group ($M = 2.74$, $SD = 0.92$) than in the non-disclosure group ($M = 2.35$, $SD = 0.90$). In contrast, purchase intention was slightly lower in the disclosure group ($M = 1.63$, $SD = 1.13$) compared to the non-disclosure group ($M = 1.72$, $SD = 1$), although the difference was minimal. Furthermore, participants who saw the high involvement product report higher involvement with the respective product ($M = 3.37$, $SD = 1.29$) than those in the low involvement condition ($M = 2.90$, $SD = 1.22$). To verify whether these ratings were statistically significant, a one-sided independent samples t-test was conducted. Results confirmed the significance ($t(134) = -2.15$, $p = .017$), meaning that participants in the low involvement condition reported significantly lower involvement than those in the high involvement condition. Therefore, the manipulation of involvement was successful.

The descriptive statistics for all measured variables separated by the control and experimental group, including sample size (N), minimum and maximum values, mean, and standard deviation, are presented in Table 3.

Table 3: Descriptive statistics of the variables

	N	Min	Max	Mean	Std. Dev.
Authenticity (non-disclosed)	71	1	4.4	2.3324	.90329
Authenticity (disclosed)	65	1	4.8	2.5908	.89053
Brand Image (non-disclosed)	71	1	5	2.3549	.89901
Brand Image (disclosed)	65	1	5	2.7415	.91701
Purchase Intention (non-disclosed)	71	1	5	1.72	1.003
Purchase Intention (disclosed)	65	1	4	1.63	1.128
Involvement (low inv. condition)	67	1	6.5	2.9030	1.21851
Involvement (high inv. condition)	69	1	6.5	3.3659	1.28665

In order to examine the hypotheses, a moderation as well as mediation analysis was performed using model one and four of the PROCESS macro by Hayes (2018, pp. 91, 238). Bootstrapping with 5,000 samples and heteroscedasticity consistent standard errors (HC4) (Cribari-Neto,

2004, p. 232) was used to calculate the confidence intervals and inferential statistics (Hayes, 2018, p. 71). This approach allowed for robust inference making regarding the indirect effect, which was considered significant when the corresponding confidence interval excluded zero (Hayes, 2018, pp. 93ff). More specifically, model one tested whether AI disclosure (X) had an effect on the perceived authenticity/brand image/purchase intention (Y) and if this was moderated by the level of involvement (W) (Hayes, 2018, p. 238). Furthermore, model four tested whether AI disclosure (X) predicted brand image respectively purchase intention (Y) and whether this relationship was fully or partially mediated by authenticity (M) (Hayes, 2018, p. 91).

3.3.2 Assumption Testing

Reliability Analysis

Cronbach's α was calculated once again for authenticity, brand image, and involvement to test the internal consistency of the measurement scales - this time with the data of the main study. Purchase intention was not tested, as it was measured by a single item. All constructs proved high reliability, with Cronbach's α ranging from .835 to .926. However, involvement showed lower reliability, with the lowest α at .675, which is still considered acceptable (Malhotra, 2020, p. 303). Detailed reliability analyses are provided in Appendix C.

Normality Testing

As parametric tests such as ANOVA and PROCESS assume normality (Hayes, 2018, p. 70), the distribution of all dependent variables was examined using the Shapiro-Wilk test, differentiating by disclosure groups. The results indicated significant deviations from normality for all variables in the non-disclosure condition ($p < .01$). In the disclosure condition, Brand image and purchase intention also showed significant departures from normality ($p = .003$; $p < .001$). Only authenticity in the disclosure condition did not significantly deviate from normality ($p = .123$), suggesting approximate normality. These findings indicate that, overall, the assumption of normality was violated for most variables. However, literature suggests not to only conduct the Shapiro-Wilk test but also generating quantile-quantile-plots (Field, 2018, p. 250). The plot for authenticity in the disclosure condition supported the statistical result, indicating an approximately normal distribution. In contrast, the Q-Q plot for brand image showed slight deviations from the reference line, suggesting mild skewness or kurtosis. The plot for purchase intention displayed clear deviations from normality. Overall, these findings confirmed the violation of normality. However, prior research demonstrated that the F-test remains robust even when normality is violated (Blanca et al., 2017, p. 554). Therefore,

ANOVA remained applicable in this study. Additionally, bootstrapping was applied in all PROCESS analyses to account for non-normality (Hayes, 2018, p. 165).

Homogeneity of Variance

Levene's test was conducted for all dependent variables to ensure that the assumption of homogeneity of variances was met and that the variance across the groups was equal. Violation of homogeneity of variance can affect the validity of the results of both the ANOVA and t-test (Field, 2018, p. 237). The results of Levene's test for authenticity, brand image, and purchase intention based on the mean, were not statistically significant ($p = .876$; $p = .542$; $p = .287$), confirming the assumption of homogeneity of variance.

Multicollinearity

Multicollinearity refers to a high correlation between two or more predictor variables in a regression model, which can lead to unstable regression coefficients and should be eliminated when using PROCESS (Bowerman et al., 2015, p. 164). To ensure the reliability of the regression results, multicollinearity was assessed by examining tolerance values and Variance Inflation Factors (VIFs) for all predictors included in the model. For the dependent variables all VIFs were below 1.008 and all tolerance values exceeded .992 testing on the independent variables disclosure and involvement. These values fall well within acceptable ranges (Bowerman et al., 2015, p. 162). Accordingly, there is no indication of problematic multicollinearity.

Linearity

Linearity is an assumption for the PROCESS analysis, however as the independent variables are binary, linearity did not need to be tested (Tabachnick & Fidell, 1989, p. 80). Linearity applies only to continuous predictors, as a straight-line model is the outcome of categorical variables.

No outliers

Boxplots were used to examine the presence of outliers in the continuous variables, separating by the disclosure condition. In the non-disclosure condition, authenticity and brand image showed one or no outlier. In both conditions, four outliers appeared for purchase intention. None were observed in both brand image and authenticity when AI was present. These observations indicate mild violations of the no-outliers assumption for Purchase Intention. Nevertheless, all values remained in the dataset, as they may represent true variation rather than measurement error. To address the potential influence of these outliers, bootstrapping was applied in all PROCESS models, ensuring robustness of the estimates even when assumptions were not fully met.

Homoscedasticity

As mentioned above, HC4 will be applied, so the assumption for homoscedasticity does not need to be tested (Cribari-Neto, 2004, p. 232).

Independence

The Durbin-Watson-Test was used to test independence (Tabachnick & Fidell, 1989, p. 133). Results showed no serious autocorrelation, as the values range from 2.174 to 2.315, suggesting that the residuals are independent and the assumption is met (Bowerman et al., 2015, p. 210).

3.3.3 Testing Hypotheses

The hypotheses testing begins with the analysis of Hypothesis 1, followed by Hypothesis 2 to examine if each hypothesis can be accepted or has to be rejected.

3.3.3.1 Hypothesis 1

To test Hypothesis 1, a moderation analysis was conducted using PROCESS Model 1 (Hayes, 2018, p. 238), with involvement as the moderator and authenticity (H1a), brand image (H1b) respectively purchase intention (H1c) as dependent variables.

For H1a, the overall model did not yield statistical significance $F(3, 132) = 1.5605, p = .2020$. The R-squared (R^2) value measures the proportion of variance in the dependent variable that is explained by the predictor variables. In this case, the R^2 value was .0337, indicating that the predictor variables explained only a small portion of the variability in perceived authenticity. The interaction between disclosure and involvement was not significant ($b = .3495, p = .261$), because the 95 % confidence interval included zero $([-.2635; .9625])$. This non-significance remained consistent at the 90 % confidence level $([-.1638; .8628])$, and in the bootstrapped intervals. These results suggest that the effect of disclosure on authenticity did not differ across involvement levels. Therefore, involvement does not act as a moderator for authenticity and H1a could not be accepted.

Table 4: ANOVA: Authenticity

	Sum of Squares	df	Mean squares	F	Sig.
Between Groups	2.265	1	2.265	2.814	.096
Within Groups	107.870	134	.805	-	-
Total	110.135	135	-	-	-

To further explore the main effect of disclosure, a separate one-way ANOVA was conducted to test whether disclosure alone, regardless of involvement, affected perceived authenticity. This analysis revealed a marginally significant effect at the 90 % confidence level ($F(1, 134)$

= 2.814, $p = .096$), meaning that disclosure had a positive effect on authenticity. In the context, while perceived authenticity depended on whether disclosure was present, this effect did not significantly vary across involvement levels.

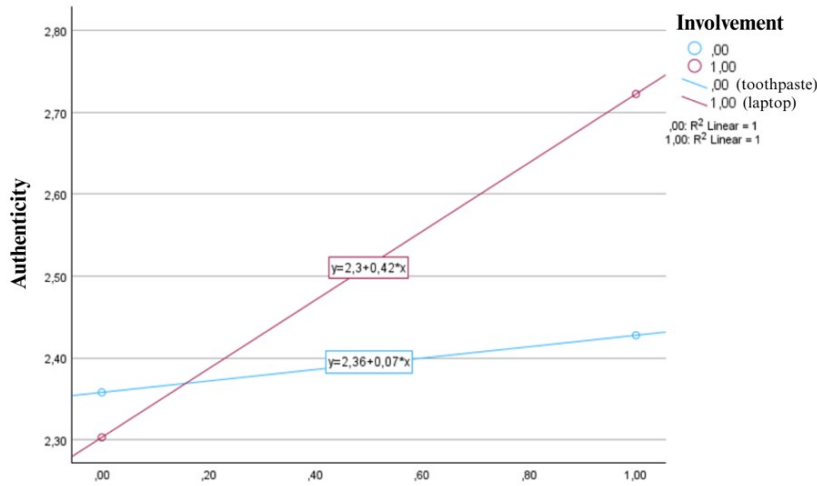


Figure 3: Graph Authenticity/Disclosure

appeared to increase perceived authenticity overall. Notably, for high-involvement products, authenticity was perceived higher when AI was disclosed.

The same analysis was conducted to test if involvement moderates the relationship between disclosure and either brand image or purchase intention. For both outcome variables, the model did not reach statistical significance (Brand Image: $F(3, 132) = 2.5527$, $p = .0583$, $R^2 = .052$; Purchase Intention: $F(3, 132) = 0.5376$, $p = .6573$, $R^2 = .010$). The confidence intervals for the interaction terms between disclosure and involvement for brand image included zero at both the 90% $[-.4368, .6120]$ and 95% levels $[-.5386, .7138]$. The same applied to purchase intention (95 % CI $[-.9079, .5636]$) as illustrated in Table 6. The bootstrapped intervals confirmed the output for both variables, showing no evidence that the effect of AI disclosure on brand image and purchase intention differed across levels of involvement.

Table 5: Moderation Analysis Disclosure x Involvement for Brand Image

	coeff	se (HC4)	t	p	CI 90 %
constant	2.3	.1366	16.8404	.0000	[2.0738, 2.5262]
Disclosure	.3276	.2217	1.4776	.1419	[-.0397, .6948]
Involvement	.1182	.2171	.5445	.5870	[-.2414, .4777]
Int_1	.0876	.3166	.2766	.7825	[-.4368, .6120]

Table 6: Test(s) of highest order unconditional interaction(s)
(Purchase Intention, 95 % CI, Interaction of Disclosure x Involvement)

	R2-chng	F(HC4)	df1	df2	p
X*W	.0016	.2142	1	132	.6443

However, the ANOVA for brand image revealed a significant difference between disclosure types ($F(1, 134) = 6.157, p = .014$), which is also confirmed by the 95 % confidence interval since it did not include zero ($[-.0784, .6948]$). The combination of the findings from both the PROCESS model and ANOVA indicates that AI disclosure influences brand image perceptions. However, this effect remained consistent regardless of involvement levels.

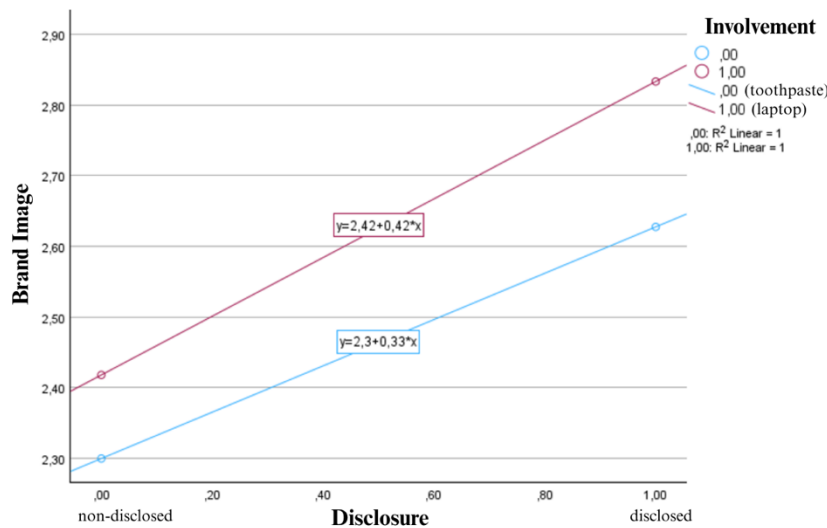


Figure 4: Graph Brand Image/Disclosure

In contrast, for purchase intention, disclosure alone had no significant effect, meaning that neither disclosure itself nor its interaction with involvement predicted participants' purchase intention.

Figure 5 presents the results of hypothesis 1, showing the main effect of AI disclosure, the main effect of involvement as well as the interaction effect between disclosure and involvement on each outcome variable. However, H1a, H1b, and H1c are rejected as no effect was significant.

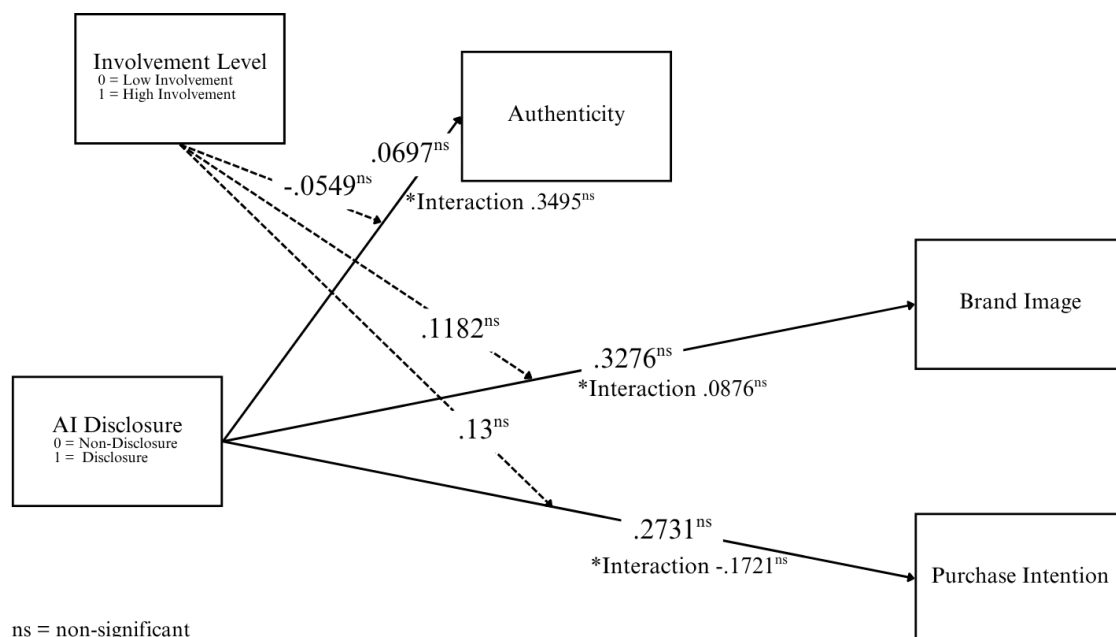


Figure 5: Results of H1

Nevertheless, visual analysis showed a generally slightly higher brand image for toothpaste than laptops especially when AI usage was disclosed (see Figure 4). These results were similar to authenticity perceptions.

3.3.3.2 Hypothesis 2

PROCESS Model 4 investigated whether authenticity mediated the relationship between AI disclosure and brand image, as well as purchase intention.

Results for Hypothesis 2a were not significant at the 95 % confidence level ($[-.0311, .4555]$). However, at the 90 % confidence level, the indirect effect showed a marginal mediation effect of authenticity on brand image ($[.0111, .4127]$; see Table 7).

Table 7: Indirect effect(s) of X on Y

	Effect	BootSE	90 % CI
Authenticity	-.2060	.1229	[.0111, .4127]

As presented in Table 8, results also showed that AI disclosure had a significant positive total effect on brand image ($b = .3866, p = .0144$), meaning that when disclosure was present, the evaluated brand image was better. When authenticity was included as a mediator, the direct effect of disclosure on brand image was reduced, but remained marginally significant on a 90 % level ($b = .1806, p = .0742$). This indicated a partial mediation of disclosure on the brand image via authenticity, supporting H2a.

Table 8: Summarized Mediation Table for Brand Image

Path	Effect	SE (HC4)	t	p	90 % CI
a: Disc → Auth	.2584	.1539	1.6785	.0956	[.0034, .5133]
b: Auth → BI	.7972	.0502	15.8894	.0000	[.7141, .8802]
c: Total effect	.3866	.1560	2.4789	.0144	[.1283, .6449]
c': Direct effect	.1806	.1004	1.7994	.0742	[.0144, .3469]
ab: Indirect effect	.2060	.1229 (BootSE)	-	-	[.0111, .4127] (BootCI)

Nevertheless, no mediation was observed at the 95 % confidence level, as both the direct and indirect confidence intervals included zero ($[-.0179, .3792]$, $[-.0311, .4555]$).

The same model was used to examine whether a mediation for purchase intention existed and revealed a significant interaction at the 90 % confidence level ($b = .2275, [.0055, .4532]$), although the effect fell short of significance at the 95 % level ($[-.0324, .5020]$). Neither the total effect ($b = .1894, p = .3044$) nor the direct effect of disclosure on purchase intention ($b = -.0381, p = .7677$) were statistically significant, as already shown in the ANOVA results. However, authenticity significantly predicted purchase intention ($b = .8806, p < .001$). These findings indicate an indirect-only mediation (Zhao et al., 2010, p. 201), supporting H2b as well.

The detailed results of the mediation analysis, including all paths, are presented in Table 9 and visually illustrated in Figure 6 (excluding the indirect effect).

Table 9: Summarized Mediation Table for Purchase Intention

Path	Effect	SE (HC4)	t	p	90 % CI
a: Disc → Auth	.2584	.1539	1.6785	.0956	[.0034, .5133]
b: Auth → PI	.8806	.0770	11.4422	.0000	[.7531, 1.0081]
c: Total effect	.1894	.1837	1.0310	.3044	[-.1149, .4936]
c': Direct effect	-.0381	.1289	-.2959	.7677	[-.2517, .1754]
ab: Indirect effect	.2275	.1364 (BootSE)	-	-	[.0055, .4532] (BootCI)

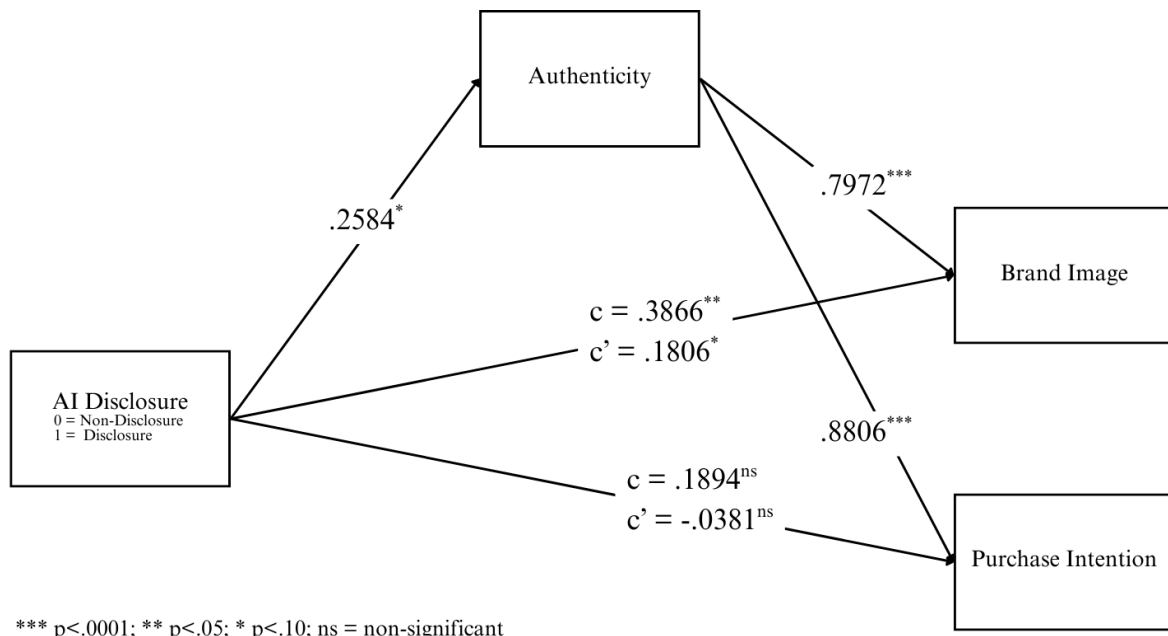


Figure 6: Results of H2

Notably, the main effect of disclosure on each outcome variable differed between the moderation model (H1) and the mediation model (H2). This difference is due to model structure: the moderation model was conditional on the inclusion of involvement and its interaction term, which adjusted the main effect of disclosure. In contrast, the mediation model included authenticity as a mediator, without controlling for any moderator, thereby reflecting disclosure's full effect on the outcome variables and leading to different total effects.

4 Discussion

In the following discussion, the key findings of this study are interpreted in light of existing literature, with particular focus on noteworthy results. Furthermore, strengths and limitations of the research design, and implications for future studies are considered.

4.1 Interpretation of Key Findings

The aim of the study was to examine the role of AI disclosure in ad creation on consumer perceptions for high involvement versus low involvement products.

Based on prior research, it has been hypothesized that AI disclosure would negatively affect authenticity, brand image, and purchase intention (Kučinskas & Survilaite, 2025, p. 727; Luo et al., 2019, p. 942). ANOVA results showed significant differences for authenticity and brand image in both groups. However, against the assumption and what literature suggested, AI disclosure significantly improved brand image. For instance, Kučinskas & Survilaite emphasize that brand image worsens when AI usage is present (2025, p. 727), and Kirkby et al. report non-significant effects of AI disclosure on brand attitude (2023, p. 1115). Regarding authenticity, prior research found a decline when AI is disclosed (Amos & Zhang, 2024, p. 7; Kučinskas & Survilaite, 2025, p. 727; Lee & Kim, 2024, p. 5). Yet, the current findings suggest marginally higher perceived authenticity when disclosure is present. Again, Kirkby et al. could not find significant differences between the groups (2023, p. 1115). The positive findings support the idea that disclosed advertisements serve as a signal of transparency and innovation, aligning with claims made by Kučinskas & Survilaite (2025, p. 727). This suggests, that brands should not fear labelling the use of AI as it does not harm consumers' perceptions. However, purchase intention does not show differences between the levels of disclosure on a 90 % confidence interval. This suggests that disclosure affects perceptual constructs rather than behavioral outcomes, aligning with the intention-behavior gap described in the literature, suggesting that improved perceptions alone are insufficient to generate stronger purchase intentions.

Hypothesis 1 additionally proposed that product involvement would moderate the effects of disclosure, based on the idea that consumers process information more carefully and intensively for high involvement compared to low involvement products (Petty & Cacioppo, 1986, pp. 4, 125). Contrary to this expectation, no significant interaction effects were found for any of the outcome variables. This suggests that AI disclosure exerts similar effects across low- and high-involvement product ads, regardless of how engaged consumers are with a product, meaning that the results are generalizable across products categories.

In summary, H1 was not supported, as neither authenticity, brand image, nor purchase intention showed interaction effects with the level of involvement. Additionally, the direction of the disclosure effect on authenticity and brand image was opposite to what was hypothesized. A possible explanation lies in the characteristics of the sample: the study primarily attracted younger and well-educated participants, a demographic typically more familiar with AI technologies. This higher degree of familiarity may have shaped their evaluations, leading to greater acceptance of AI disclosure compared to what might be expected in a more diverse population.

Results showed partial support for H2a: authenticity marginally mediated the effect of AI disclosure on brand image on a 90 % confidence level. Similarly, the effect of disclosure on purchase intention was also mediated by authenticity at the same confidence level. Indirect effects for both outcome variables were positive, meaning that disclosure increased authenticity (path a), and that authenticity, in turn, increased both brand image and purchase intention (path b) (Hayes, 2018, p. 84). Hence, higher perceived authenticity leads to a more favorable brand image and greater purchase intention. This aligns with previous research showing that authenticity mediates the effects of AI disclosure on brand image (H2a) and purchase intention (H2b) (Kirkby et al., 2023, p. 1115; Kučinskas & Survilaite, 2025, p. 727). Nevertheless, Kučinskas & Survilaite found that authenticity negatively mediated the relationship between AI disclosure and consumers evaluations (2025, p. 727), indicating an opposite effect direction compared to the findings in the present study.

The total and direct effects of disclosure on brand image were significant at the 90 % confidence level, supporting a partial mediation model. Therefore, authenticity partly explains the positive effect of AI disclosure on brand image. In contrast, for purchase intention, neither the total nor the direct effect yielded significant results, indicating an indirect-only mediation. That means that disclosure affects purchase intention solely through its influence on perceived authenticity, rather than having a direct effect itself (Zhao et al., 2010, p. 201). However, findings in the literature also reported direct effects of disclosure on purchase intention respectively purchase rates (Kučinskas & Survilaite, 2025, p. 732; Luo et al., 2019, p. 942).

4.2 Noteworthy Results

Several key findings emerged from the analysis, holding unexpected results in light of previous research.

Most notably, AI disclosure led to higher perceived authenticity, albeit marginally, and significantly improved brand image. While literature suggested that consumers evaluations

actually worsen when the AI label is present. Similarly, while it was hypothesized that disclosure would reduce purchase intention, the results showed no significant difference between present and absent disclosure.

Interestingly, the level of involvement also did not moderate any of the outcome variables. Although no interaction was statistically significant, visual analysis revealed a notable pattern: authenticity was generally low for the low involvement product, but slightly higher for high involvement products when AI was disclosed.

As anticipated, authenticity acted as a mediator for both brand image and purchase intention as the indirect effect was significant. Notably, the direct effect of disclosure on purchase intention was not significant, suggesting that authenticity serves as a key mediator.

It is worth noting that most results did not reach statistical significance at the 95 % confidence level, only at the 90 % level. The only exception was the positive main effect of disclosure on brand image indicated significance at the 95 % confidence level, highlight this as the most robust and noteworthy finding of the study.

4.3 Strengths and Limitations

This study presents several strengths but also faces certain limitations that should be considered when interpreting the results.

One of the key strengths lies in the experimental design, which supported a high level of both internal and external validity. External validity was enhanced by randomly allocating participants to conditions, which minimized systematic differences between groups. Internal validity benefited from the use of hypothetical brands, which prevented participants from holding pre-existing opinions or associations. However, it may reduce external validity, as consumer responses to real, familiar brands may differ - especially considering prior research that suggests AI disclosure is perceived more negatively when associated with well-known brands (Kučinskas & Survilaite, 2025, p. 727).

To ensure construct validity, all key variables were measured using established and reliable scales drawn from existing literature. However, purchase intention scale is not a thoroughly reliable indication for subsequent behavior, known as the intention-behavior-gap (Conner & Norman, 2022, p. 2; Morwitz & Munz, 2021, p. 31; Wong & Sheth, 1985, p. 378).

The application of PROCESS models with heteroscedasticity-consistent standard errors (HC4) and bootstrapped confidence intervals strengthened the robustness of the results, especially given the non-normality of some variables (Cribari-Neto, 2004, p. 232).

The visual consistency of both advertisements further contributed to internal validity by reducing the risk of design-based bias. However, this simplicity of the ads may have contributed to the generally low ratings across variables, as the advertisements were not unique or attention-grabbing.

Additionally, the chosen product categories may carry specific connotations that could influence responses. While toothpaste is likely to be viewed as a functional product, a laptop is associated with technology - potentially leading participants to view AI-generated advertising for laptops as more acceptable. This may partly explain the non-significant moderating effect of involvement on the relationship between AI disclosure and consumer evaluations.

Another limitation concerns the sample composition. The sample group was primarily young, well-educated individuals, mostly from Generation Z. This limits the generalizability of the results to a broader consumer population as younger and more educated individuals may have greater familiarity or openness to AI technologies, which could influence their perceptions in a way not representative to the broader population. Additionally, factors such as participants' technological affinity or general attitudes toward Artificial Intelligence may have influenced responses but were not directly measured or controlled for.

However, participants were asked whether they perceived the advertisement as AI-generated. In the control group - where the AI label was not present - 53.5 % still identified the ad as AI-generated. This may have affected authenticity ratings due to a perceived lack of transparency, potentially negatively biasing results in the control group and confounding the manipulation.

Another limitation concerns the statistical significance level of the results. Most findings reached significance only at the 90 % confidence level, indicating a marginal level of certainty. The only results that were statistically significant at the 95 % confidence level were the indirect effect of authenticity on purchase intention as well as the main effect of disclosure on brand image, suggesting more robust relationships in these cases. However, the overall reliance on a lower confidence level may limit the strength and generalizability of the conclusions drawn from the data.

4.4 Future Research

The findings and limitations of this study suggest several directions for future research.

With the upcoming mandatory disclosure of AI-generated content from August 2026 onward, AI labels will become more important in marketing contexts. Future studies should investigate how this regulatory shift affects consumer expectations, trust, and ad effectiveness. Moreover,

research should test which disclosure framing strategies - such as the length and wording of disclosure or highlighting human involvement - lead to the most favorable outcomes.

Second, individual characteristics may play a crucial role in the evaluation of AI usage in advertising. Liang discovered that women, older people, less educated individuals and those with lower income are more likely to experience fear toward artificial intelligence (2017, p. 383), suggesting that demographic variables influence AI acceptance. However, these results may have shifted in recent years as AI became more present in daily life. Additionally, technical affinity or a person's degree of innovativeness may be relevant characteristics that could be examined.

Third, the industry context in which AI is disclosed may play a crucial role. Certain sectors such as fashion, healthcare, technology or finance evoke different expectations regarding ethics, creativity and personal risk. Therefore, usage of AI in advertising context may be valued differently across industries, which could be investigated in future studies to determine if AI is more accepted or problematic in specific sectors. Furthermore, privacy concerns may also be relevant in the context of personalized advertising through artificial intelligence. Although AI might be a useful tool to personalize ads effectively, it might also influence consumer behavior if they know their personal data was analyzed with AI for the advertising.

Fourth, future studies should move beyond laboratory, self-reported evaluations and instead observe actual behavior. While this study investigated purchase intention, field experiments could test if the real purchase differs from the claimed one and whether there is an intention-behavior gap. Additionally, future research could examine how AI-generated advertisements are perceived when real brands are used. Kučinskas & Survilaite (2025, p. 734) found that consumers evaluate AI-created ads more negatively when they are familiar with the brand. However, it remains unclear whether this varies based on consumers' preexisting views of the brand.

5 Conclusion

This thesis set out to explore how consumers evaluate AI-generated advertising, particularly when the use of AI is disclosed and presented in the context of products with differing levels of involvement. By experimentally testing the effects of AI disclosure on perceived authenticity, brand image, and purchase intention, the study offers valuable insights into how transparency about AI usage influences consumer perception in advertising.

Contrary to much of the existing literature, the findings reveal that AI disclosure can positively shape certain consumer evaluations. Specifically, this study found that transparency can actually enhance both authenticity and brand image. Especially, when consumers perceived the ad as authentic, they were more likely to evaluate the brand image positively and express higher purchase intentions. However, AI disclosure did not influence purchase intention directly, instead, the effect occurred indirectly through authenticity, highlighting its mediating role.

The study also analyzed whether these effects varied depending on the level of involvement, based on the Elaboration Likelihood Model (Petty & Cacioppo, 1986). Theoretically, high involvement is associated with increased cognitive processing and was therefore expected to respond more critically to AI disclosure. In contrast, low involvement conditions typically require a lower cognitive effort, suggesting that disclosure would have less impact. Empirically, however, consumer evaluations did not significantly differ across involvement levels, rejecting hypothesis 1. Nonetheless, visual analyses suggested that for high involvement products, authenticity and brand image were perceived as higher when AI was disclosed, indicating that involvement plays a smaller role than previously assumed, or that its effects are more context dependent.

Overall, this study contributes to a better understanding of AI in advertising. It suggests that AI disclosure does not inherently harm consumer perceptions and may even enhance it, as it is a sign of transparency and innovation. However, these results are particularly relevant for brands targeting a young, well-educated consumers, as this demographic was most strongly represented in the sample.

Especially, in light of the AI Act, this research offers timely and practical insights. It underlines the importance of how AI is communicated and suggest that strategic communication of disclosure can significantly impact brand evaluations. Companies should therefore consider how they integrate AI use in their advertising strategies to maintain consumer trust and strengthen brand positioning.

Bibliography

- Ajzen, I. (1991). The theory of planned behavior. *Organizational Behavior and Human Decision Processes*, 50(2), 179–211. [https://doi.org/10.1016/0749-5978\(91\)90020-T](https://doi.org/10.1016/0749-5978(91)90020-T)
- Ajzen, I., & Fishbein, M. (1980). Understanding attitudes and predicting social behavior. Prentice-Hall.
- Akbari, M. (2015). Different Impacts of Advertising Appeals on Advertising Attitude for High and Low Involvement Products. *Global Business Review*, 16(3), 478–493. <https://doi.org/10.1177/0972150915569936>
- Amos, C., & Zhang, L. (2024). Consumer reactions to perceived undisclosed ChatGPT usage in an online review context. *Telematics and Informatics*, 93, 102163. <https://doi.org/10.1016/j.tele.2024.102163>
- Arango, L., Singaraju, S. P., & Niininen, O. (2023). Consumer Responses to AI-Generated Charitable Giving Ads. *Journal of Advertising*, 52(4), 486–503. <https://doi.org/10.1080/00913367.2023.2183285>
- Baati, O., & Akrouit, F. (2024). How long does purchase intention exist? “Buying intention survival”: a new insight into forecasting purchase intention abandonment. *The Journal of Consumer Marketing*, 41(7), 734–750. <https://doi.org/10.1108/JCM-12-2023-6474>
- 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>
- Barari, M., Casper Ferm, L.-E., Quach, S., Thaichon, P., & Ngo, L. (2024). The dark side of artificial intelligence in marketing: meta-analytics review. *Marketing Intelligence & Planning*, 42(7), 1234–1256. <https://doi.org/10.1108/MIP-09-2023-0494>
- Beverland, M. B., & Farrelly, F. J. (2010). The Quest for Authenticity in Consumption: Consumers’ Purposive Choice of Authentic Cues to Shape Experienced Outcomes. *Journal of Consumer Research*, 36(5), 838–856. <https://doi.org/10.1086/615047>
- Bowerman, B. L., O’Connell, R. T., & Murphree, E. S. (2015). Regression analysis: Unified concepts, practical applications, and computer implementation (First edition). Business Expert Press.
- Bünthe, C. (2023). *Studie: Künstliche Intelligenz – die Zukunft des Marketings 2023*. SRH Berlin University of Applied Sciences. https://www.srh-university.de/fileadmin/Hochschule_Berlin/B.A._IBWL_Marketing/Studie_KI_im_Marketing_2023_final_1_.pdf

- Campbell, C., Plangger, K., Sands, S., & Kietzmann, J. (2022). Preparing for an Era of Deepfakes and AI-Generated Ads: A Framework for Understanding Responses to Manipulated Advertising. *Journal of Advertising*, 51(1), 22–38.
<https://doi.org/10.1080/00913367.2021.1909515>
- Celuch, K. G., & Taylor, S. A. (1999). Involvement with Services: An Empirical Replication and Extension of Zaichkowsky's Personal Involvement Inventory. *The Journal of Consumer Satisfaction, Dissatisfaction and Complaining Behavior*, 12, 109–122.
- Conner, M., & Norman, P. (2022). Understanding the intention-behavior gap: The role of intention strength. *Frontiers in Psychology*, 13, 923464.
<https://doi.org/10.3389/fpsyg.2022.923464>
- Cribari-Neto, F. (2004). Asymptotic inference under heteroskedasticity of unknown form. *Computational Statistics & Data Analysis*, 45(2), 215–233.
[https://doi.org/10.1016/S0167-9473\(02\)00366-3](https://doi.org/10.1016/S0167-9473(02)00366-3)
- de Vries, A. (2023). The growing energy footprint of artificial intelligence. *Joule*, 7(10), 2191–2194. <https://doi.org/10.1016/j.joule.2023.09.004>
- Dobni, D., & Zinkhan, G. M. (1990). In Search of Brand Image: A Foundation Analysis. *Advances in Consumer Research*, 17, 110–119.
- Duolingo @Duolingo (2025, April). ☞ Below is an all-hands email from our CEO, Luis von Ahn - we are going to be AI-first. [LinkedIn post]. LinkedIn.
https://www.linkedin.com/posts/duolingo_below-is-an-all-hands-email-from-our-activity-7322560534824865792-19vh?utm_source=share&utm_medium=member_desktop&rcm=ACoAADJDcb4BpN86xjUwS25fRB1cjUrbgrcgXOE
- Dutton, D. (2003). Authenticity in art. In J. Levinson (Ed.), *The Oxford handbook of aesthetics* (S. 258–274). Oxford University Press.
<https://citeseerx.ist.psu.edu/document?repid=rep1&type=pdf&doi=69c93a55f274c7666fb51da32d0c1f2268b5e67d>
- Eisend, M. (2013). The Moderating Influence of Involvement on Two-Sided Advertising Effects. *Psychology & marketing*, 30(7), 566–575.
- European Parliament. (2025, February 19). EU AI Act: First Regulation on Artificial Intelligence. <https://www.europarl.europa.eu/topics/en/article/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence>

- Feuerriegel, S., Hartmann, J., Janiesch, C., & Zschech, P. (2024). Generative AI. *Business & Information Systems Engineering*, 66(1), 111–126. <https://doi.org/10.1007/s12599-023-00834-7>
- Geuens, M., & De Pelsmacker, P. (2017). Planning and Conducting Experimental Advertising Research and Questionnaire Design. *Journal of Advertising*, 46(1), 83–100. <https://doi.org/10.1080/00913367.2016.1225233>
- Grigsby, J. L., Michelsen, M., & Zamudio, C. (2025). Service ads in the era of generative AI: Disclosures, trust, and intangibility. *Journal of Retailing and Consumer Services*, 84, 104231. <https://doi.org/10.1016/j.jretconser.2025.104231>
- Haleem, A., Javaid, M., Asim Qadri, M., Pratap Singh, R., & Suman, R. (2022). Artificial intelligence (AI) applications for marketing: A literature-based study. *International Journal of Intelligent Networks*, 3, 119–132. <https://doi.org/10.1016/j.ijin.2022.08.005>
- Hayes, A. F. (2018). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach* (Second edition). Guilford Press.
- Huh, J., Nelson, M. R., & Russell, C. A. (2023). ChatGPT, AI Advertising, and Advertising Research and Education. *Journal of Advertising*, 52(4), 477–482. <https://doi.org/10.1080/00913367.2023.2227013>
- Jain, M. (2019). A study on consumer behavior-decision making under high and low involvement situations. *International Journal of Research and Analytical Reviews*, 1(6), 943–947.
- Karpinska-Krakiowiak, M., & Eisend, M. (2024). Realistic Portrayals of Untrue Information: The Effects of Deepfaked Ads and Different Types of Disclosures. *Journal of Advertising*, 1–11. <https://doi.org/10.1080/00913367.2024.2306415>
- Kirkby, A., Baumgarth, C., & Henseler, J. (2023). To disclose or not disclose, is no longer the question – effect of AI-disclosed brand voice on brand authenticity and attitude. *Journal of Product & Brand Management*, 32(7), 1108–1122. <https://doi.org/10.1108/JPBM-02-2022-3864>
- Kučinskas, G., & Survilaitė, E. (2025). Revealing AI Involvement in Ad Creation: Effects on Authenticity, Brand Perceptions and Consumer Intentions. *Journal of Information Systems Engineering and Management*, 10(16s), 727–740. <https://doi.org/10.52783/jisem.v10i16s.2659>
- Kuß, A., Wildner, R., & Kreis, H. (2014). *Marktforschung: Grundlagen der Datenerhebung und Datenanalyse* (5th ed). Retrieved 25th February 2025, from

https://usearch.uaccess.univie.ac.at/primo-explore/fulldisplay/UWI_alma51350652630003332/UWI

- Lee, G., & Kim, H.-Y. (2024). Human vs. AI: The battle for authenticity in fashion design and consumer response. *Journal of Retailing and Consumer Services*, 77, 103690. <https://doi.org/10.1016/j.jretconser.2023.103690>
- Lehman, D. W., O'Connor, K., Kovács, B., & Newman, G. E. (2019). Authenticity. *Academy of Management Annals*, 13(1), 1–42. <https://doi.org/10.5465/annals.2017.0047>
- Lim, S., & Schmälzle, R. (2024). The effect of source disclosure on evaluation of AI-generated messages. *Computers in Human Behavior: Artificial Humans*, 2(1), 100058. <https://doi.org/10.1016/j.chbah.2024.100058>
- Liang, Y., & Lee, S. A. (2017). Fear of Autonomous Robots and Artificial Intelligence: Evidence from National Representative Data with Probability Sampling. *International Journal of Social Robotics*, 9(3), 379–384. <https://doi.org/10.1007/s12369-017-0401-3>
- Luo, X., Tong, S., Fang, Z., & Qu, Z. (2019). Frontiers: Machines vs. Humans: The Impact of Artificial Intelligence Chatbot Disclosure on Customer Purchases. *Marketing Science*. <https://doi.org/10.1287/mksc.2019.1192>
- Malhotra, N. K. (2020). *Marketing research: An applied orientation* (7th edition, global edition). Pearson.
- Mehta, P., Jebarajakirthy, C., Maseeh, H. I., Anubha, A., Saha, R., & Dhanda, K. (2022). Artificial intelligence in marketing: A meta-analytic review. *Psychology & Marketing*, 39(11), 2013–2038. <https://doi.org/10.1002/mar.21716>
- Morwitz, V. G., & Munz, K. P. (2021). Intentions. *Consumer Psychology Review*, 4(1), 26–41. <https://doi.org/10.1002/arcp.1061>
- Patterson, D., Gonzalez, J., Le, Q., Liang, C., Munguia, L.-M., Rothchild, D., So, D., Texier, M., & Dean, J. (2021). *Carbon Emissions and Large Neural Network Training*. <https://doi.org/10.48550/ARXIV.2104.10350>
- Petty, R. E., & Cacioppo, J. T. (1984). The Elaboration Likelihood Model of Persuasion. *Advances in consumer research*, 11, 673–675.
- Petty, R. E., & Cacioppo, J. T. (1986). *Communication and Persuasion*. Springer New York. <https://doi.org/10.1007/978-1-4612-4964-1>
- Petty, R. E., Cacioppo, J. T., & Schumann, D. (1983). Central and Peripheral Routes to Advertising Effectiveness: The Moderating Role of Involvement. *Journal of Consumer Research*, 10(2), 135. <https://doi.org/10.1086/208954>

- Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 Laying down Harmonised Rules on Artificial Intelligence and Amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (eu) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act) (2024). https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=OJ:L_202401689&qid=1742464131033
- Spears, N., & Singh, S. N. (2004). Measuring Attitude toward the Brand and Purchase Intentions. *Journal of Current Issues & Research in Advertising*, 26(2), 53–66. <https://doi.org/10.1080/10641734.2004.10505164>
- Wong, J. K., & Sheth, J. N. (1985). Explaining Intention-Behavior Discrepancy - A Paradigm. *Advances in consumer research*, 12(1), 378–384.
- Wu, L., Dodoo, N. A., & Wen, T. J. (2025). Disclosing AI's Involvement in Advertising to Consumers: A Task-Dependent Perspective. *Journal of Advertising*, 54(1), 20–38. <https://doi.org/10.1080/00913367.2024.2309929>
- Xia, Q., He, X., & Gong, S. (2025). AI-disclosure and brand modernity: The role of “AI=novelty” lay belief. *Journal of Product & Brand Management*. <https://doi.org/10.1108/JPBM-11-2024-5575>
- Zaichkowsky, J. L. (1985). Measuring the Involvement Construct. *Journal of Consumer Research*, 12(2), 341–352.
- Zaichkowsky, J. L. (1986). Conceptualizing Involvement. *Journal of Advertising*, 15(2), 4–34. <https://doi.org/10.1080/00913367.1986.10672999>
- Zaichkowsky, J. L. (1994). The Personal Involvement Inventory: Reduction, Revision, and Application to Advertising. *Journal of Advertising*, 23(4), 59–70. <https://doi.org/10.1080/00913367.1943.10673459>
- Zhao, X., Lynch, J. G., & Chen, Q. (2010). Reconsidering Baron and Kenny: Myths and Truths about Mediation Analysis. *Journal of Consumer Research*, 37(2), 197–206. <https://doi.org/10.1086/651257>

Appendix

Appendix A: Survey

Introduction

Welcome!

Thank you for taking the time to participate in this survey.

This study is part of my **master's thesis** and aims to gain **deeper insights into how brands are perceived** in relation to their marketing strategies.

There are **no right or wrong answers** — I'm simply interested in your **honest opinions**.

The survey won't take **longer than 5 minutes**.

Please note that **all responses are completely anonymous** and will be used solely for academic research purposes.

If any questions occur, please contact: Jennifer Böhm, a12351848@unet.univie.ac.at

☐ I confirm that I am participating in this survey voluntarily and understand that my responses will remain anonymous and confidential.

Demographics

1. What is your gender?

- ☐ female
☐ male
☐ other: _____

2. How old are you?

- ☐ younger than 14 years old
☐ 14 to 29 years old
☐ 30 to 44 years old
☐ 45 to 68 years old
☐ 69 years or older

3. What is the highest level of education you have completed?

- ☐ Still in school
☐ Finished school with no qualifications
☐ (General) Certificate of Secondary Education (Haupt-/ Realschulabschluss)
☐ A-Levels (Matura/Abitur)
☐ Completion of specialized secondary school (Fachschulabschluss)
☐ Bachelor's degree

☐ Master's degree

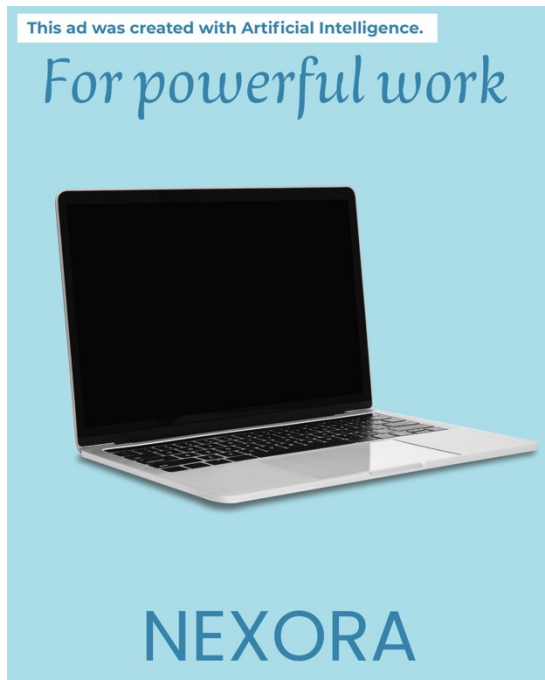
☐ Doctorate or PhD

☐ Other: _____

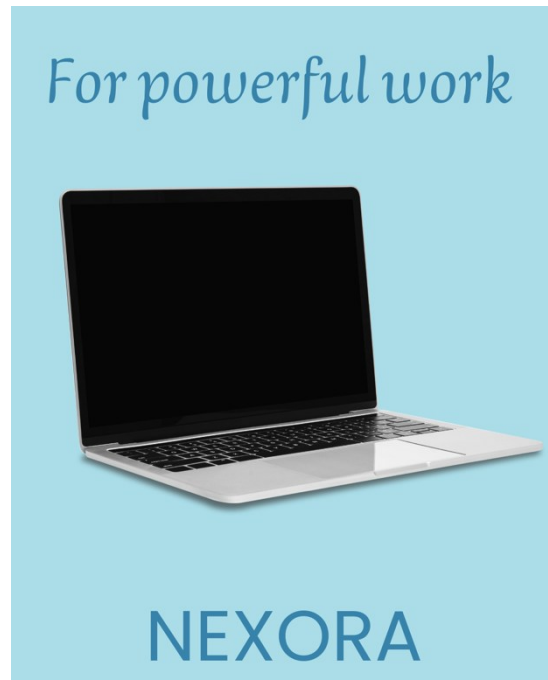
Stimuli

Please have a look at the advertisement and answer the questions below.

Stimuli 1



Stimuli 2



Stimuli 3



Stimuli 4



4. This brand seems...

Ingenuine	☐	☐	☐	☐	☐	Genuine
Not authentic	☐	☐	☐	☐	☐	Authentic
Not original	☐	☐	☐	☐	☐	Original
Unfaithful	☐	☐	☐	☐	☐	Faithful
Unnatural	☐	☐	☐	☐	☐	Natural

5. I find the brand...

Unappealing	☐	☐	☐	☐	☐	Appealing
Bad	☐	☐	☐	☐	☐	Good
Unpleasant	☐	☐	☐	☐	☐	Pleasant
Unfavourable	☐	☐	☐	☐	☐	Favourable
Unlikeable	☐	☐	☐	☐	☐	Likeable

6. How likely are you to purchase the product after seeing the ad?

Probably not buy	☐	☐	☐	☐	☐	Probably buy
------------------	--------------------------	--------------------------	--------------------------	--------------------------	--------------------------	--------------

7. Which product did you see in the ad on the previous page?

- ☐ Laptop
- ☐ Car
- ☐ Toothpaste
- ☐ Backpack
- ☐ None of these

Manipulation Check

8. How do you perceive the product shown in the advertisement?

Please place your check mark accordingly.

Unexciting	☐	☐	☐	☐	☐	☐	☐	Exciting
Appealing	☐	☐	☐	☐	☐	☐	☐	Unappealing
Means a lot to me	☐	☐	☐	☐	☐	☐	☐	Means nothing to me
Valuable	☐	☐	☐	☐	☐	☐	☐	Worthless

9. This ad seems to have been created with AI.

- ☐ Yes
- ☐ No
- ☐ Can't say

Completion

Thank you for completing the survey!

I would like to thank you very much for helping me.

Your answers were transmitted, you may close the browser tab now.

If any questions occurred, please contact a12351848@unet.univie.ac.at.

Appendix B: Pre-Test

Cronbach's Alpha (Authenticity)

Reliability Statistics

Cronbach's Alpha	N of Items
.836	5

Cronbach's Alpha (Brand Image)

Reliability Statistics

Cronbach's Alpha	N of Items
.936	5

Cronbach's Alpha (Purchase Intention)

Reliability Statistics

Cronbach's Alpha	N of Items
.950	5

Cronbach's Alpha (Involvement for toothpaste)

Reliability Statistics

Cronbach's Alpha	N of Items
.802	4

Cronbach's Alpha (Involvement for bubble bath)

Reliability Statistics

Cronbach's Alpha	N of Items
.921	4

Cronbach's Alpha (Involvement for laptop)

Reliability Statistics

Cronbach's Alpha	N of Items
.994	4

Cronbach's Alpha (Involvement for car)

Reliability Statistics

Cronbach's Alpha	N of Items
.974	4

Paired Samples T- Test (Involvement)

Paired Samples Statistics

		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	Toothpaste	3.4773	11	1.33910	.40375
	Laptop	5	11	1.81315	.54668

Paired Samples Correlation

		N	Correlation	Sig. (1-tailed p)	Sig. (2-tailed p)
Pair 1	Toothpaste & Laptop	11	.270	.211	.421

Paired Samples Test

	Mean	Std. Deviation	Std. Error Mean	Lower 95 % CI	Upper 95 % CI	t	df	Sig. (1-tailed)	Sig. (2-tailed)
Toothpaste - Laptop	-1.52273	1.94118	.58529	-2.82638	-.21862	-2.602	10	.013	.026

Appendix C: Participants Demographics

Gender	Count	Percentage
Female	93	68.4 %
Male	42	30.9 %
Other	1	0.7 %

Age	Count	Percentage
Generation Z (14-29)	98	72.1 %
Generation Y (30-44)	14	10.3 %
Generation X (45-68)	24	17.6 %

Education	Count	Percentage
Still in school	1	0.7 %
Finished school with no qualifications	0	0 %
(General) Certificate of Secondary Education (Haupt-/ Realschulabschluss)	2	1.5 %
A-Levels (Matura/Abitur)	28	20.6 %
Completion of specialized secondary school (Fachschulabschluss)	8	5.9 %
Bachelor's degree	56	41.2 %
Master's degree	35	25.7 %
Doctorate or PhD	2	1.5 %
Other	4	2.9 %

Appendix D: Main Study Results

Cronbach's Alpha (Brand Image)

Item-Skala-Statistiken					
	Skalenmittelwert, wenn Item weggelassen	Skalenvarianz, wenn Item weggelassen	Korrigierte Item-Skala-Korrelation	Quadrierte multiple Korrelation	Cronbachs Alpha, wenn Item weggelassen
Brand Image: Unappealing/Appealing	10,40	13,576	,818	,673	,925
Brand Image: Bad/Good	10,04	14,413	,811	,668	,926
Brand Image: Unpleasant/Pleasant	10,11	13,669	,867	,752	,915
Brand Image: Unfavourable/Favourable	10,16	13,959	,824	,689	,923
Brand Image: Unlikeable/Likeable	10,07	13,891	,835	,709	,921

Cronbach's Alpha (Authenticity)

Item-Skala-Statistiken					
	Skalenmittelwert, wenn Item weggelassen	Skalenvarianz, wenn Item weggelassen	Korrigierte Item-Skala-Korrelation	Quadrierte multiple Korrelation	Cronbachs Alpha, wenn Item weggelassen
Authentizität: Ingenuine/Genuine	9,68	13,406	,739	,565	,838
Authentizität: Not authentic/Authentic	9,76	13,133	,751	,592	,835
Authentizität: Not original/Original	10,18	13,169	,674	,473	,855
Authentizität: Unfaithful/Faithful	9,74	14,118	,679	,473	,853
Authentizität: Unnatural/Natural	9,76	13,500	,672	,465	,854

Cronbach's Alpha (Involvement)

Item-Skala-Statistiken					
	Skalenmittelwert, wenn Item weggelassen	Skalenvarianz, wenn Item weggelassen	Korrigierte Item-Skala-Korrelation	Quadrierte multiple Korrelation	Cronbachs Alpha, wenn Item weggelassen
Involvement: unexciting/exciting	10,11	17,921	,490	,246	,774
Involvement: appealing/unappealing	9,15	15,104	,568	,333	,740
Involvement: means a lot to me/means nothing to me	9,68	14,546	,613	,407	,715
Involvement: valuable/worthless	8,72	14,751	,690	,483	,675

Independent Samples T-Test (Involvement)

Group Statistics

	Involvement	N	Mean	Std. Deviation	Std. Error Mean
Involvement	0 (Low Involvement)	67	2.9030	1.21851	.14886
	1 (High Involvement)	69	3.3659	1.28665	.15489

Independent Samples Test

		Levene's-Test for Equality of Variances		t	df	Significance		t-test for Equality of Means		95 % Confidence Interval of the Difference	
		F	Sig.			1-tailed p	2-tailed p	Mean Difference	Std. Error Difference	Lower	Upper
Involvement	Equal variances assumed	.007	.931	-2.153	134	.017	.033	-.46296	.21501	-.88820	-.03771
	Equal variances not assumed			-2.155	133.918	.016	.033	-.46296	.21483	-.88786	-.03805

Independent Samples Effect Size

		Standardizer	Point Estimate	95 % Confidence Interval	
				Lower	Upper
Involvement	Cohen's d	1.25355	-.369	-.708	-.030
	Hedges' correction	1.26062	-.367	-.704	-.029
	Glass' Delta	1.28665	-.360	-.700	-.017

Assumptions

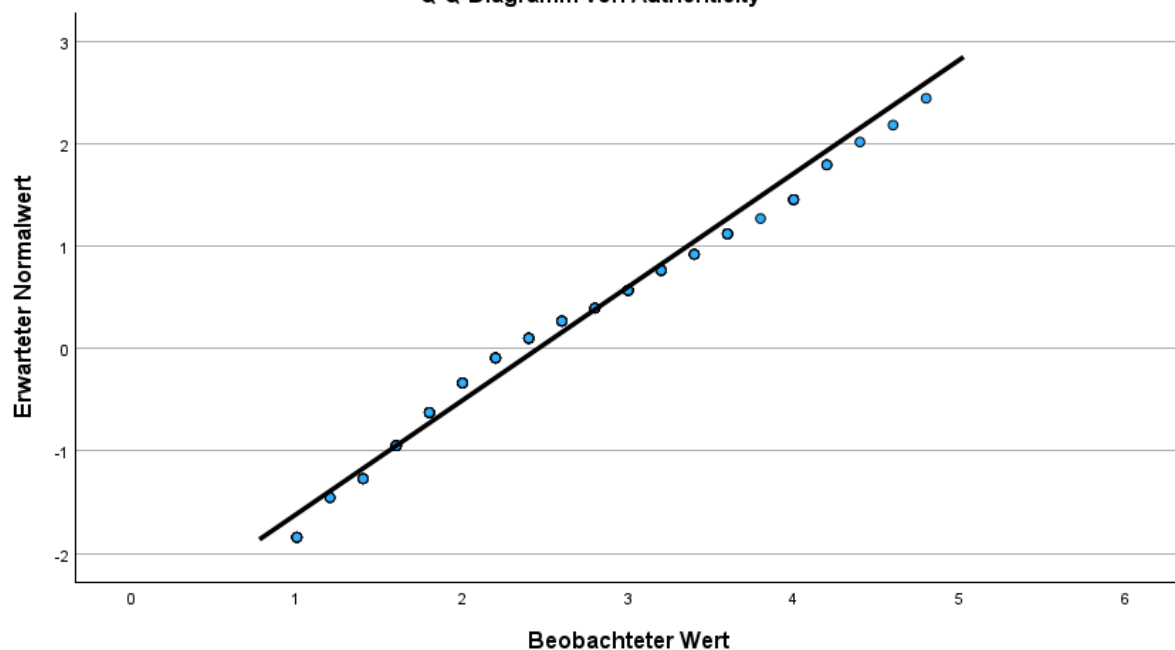
Normality Testing

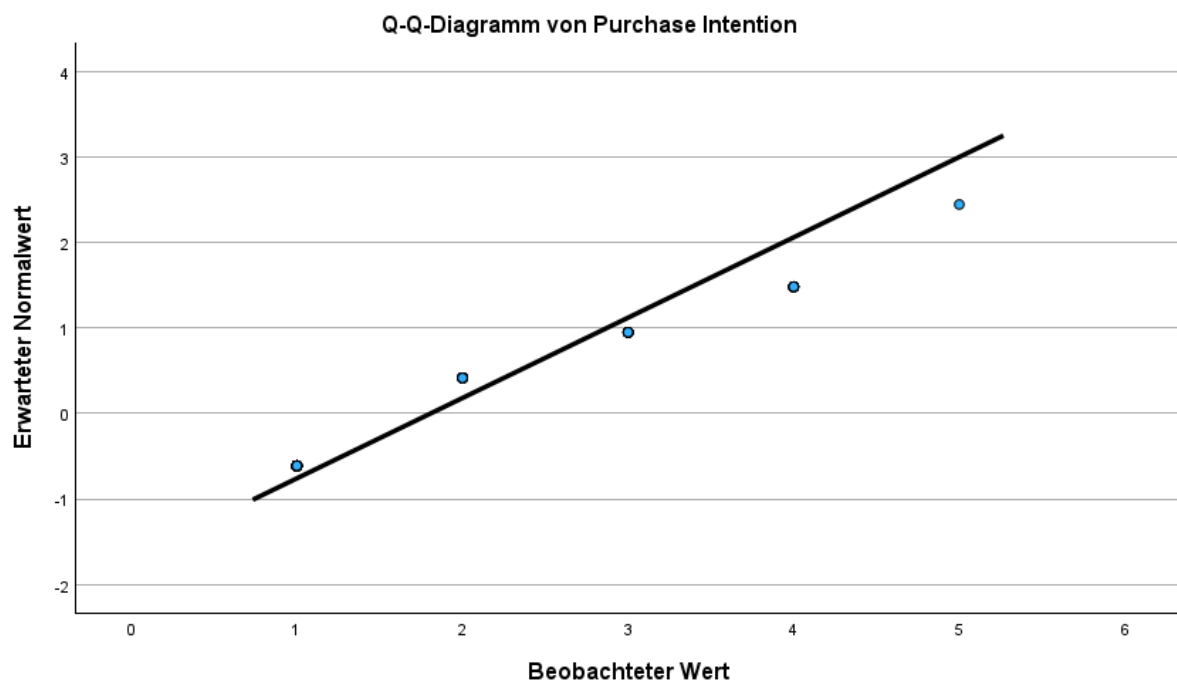
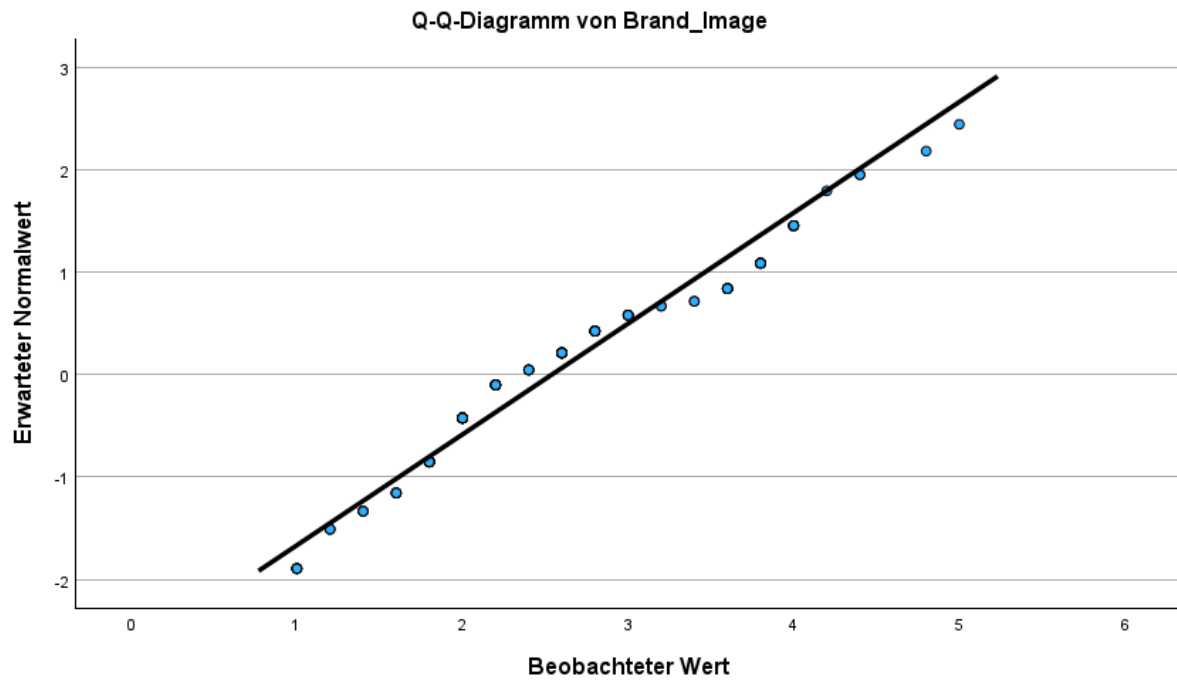
Tests auf Normalverteilung

		Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Disclosure	Statistik	df	Signifikanz	Statistik	df	Signifikanz
Authenticity	,00	,165	71	<,001	,943	71	,003
	1,00	,116	65	,031	,971	65	,123
Brand_Image	,00	,160	71	<,001	,943	71	,003
	1,00	,129	65	,009	,940	65	,003
Purchase Intention	,00	,327	71	<,001	,731	71	<,001
	1,00	,297	65	<,001	,745	65	<,001

a. Signifikanzkorrektur nach Lilliefors

Q-Q-Diagramm von Authenticity





Homogeneity of Variances Testing

Tests der Varianzhomogenität

		Levene-Statistik	df1	df2	Sig.
Authenticity	Basiert auf dem Mittelwert	,024	1	134	,876
	Basiert auf dem Median	,000	1	134	,995
	Basierend auf dem Median und mit angepaßten df	,000	1	131,374	,995
	Basiert auf dem getrimmten Mittel	,018	1	134	,895

Tests der Varianzhomogenität

		Levene-Statistik	df1	df2	Sig.
Brand_Image	Basiert auf dem Mittelwert	,374	1	134	,542
	Basiert auf dem Median	,479	1	134	,490
	Basierend auf dem Median und mit angepaßten df	,479	1	132,374	,490
	Basiert auf dem getrimmten Mittel	,425	1	134	,515

Tests der Varianzhomogenität

		Levene-Statistik	df1	df2	Sig.
Purchase Intention	Basiert auf dem Mittelwert	1,145	1	134	,287
	Basiert auf dem Median	1,074	1	134	,302
	Basierend auf dem Median und mit angepaßten df	1,074	1	132,164	,302
	Basiert auf dem getrimmten Mittel	1,332	1	134	,250

Multicollinearity Testing

Koeffizienten^a

Modell		Nicht standardisierte Koeffizienten		Standardisierte Koeffizienten	T	Sig.	Kollinearitätsstatistik	
		Regressionskoeffizient B	Std.-Fehler	Beta			Toleranz	VIF
1	(Konstante)	2,281	,130		17,556	<,001		
	Disclosure	,372	,156	,202	2,381	,019	,992	1,008
	Involvement	,160	,156	,087	1,023	,308	,992	1,008

a. Abhängige Variable: Brand_Image

Koeffizienten^a

Modell		Nicht standardisierte Koeffizienten		Standardisierte Koeffizienten	T	Sig.	Kollinearitätsstatistik	
		Regressionskoeffizient B	Std.-Fehler	Beta			Toleranz	VIF
1	(Konstante)	1,696	,153		11,093	<,001		
	Disclosure	,185	,184	,087	1,006	,316	,992	1,008
	Involvement	,048	,184	,023	,261	,795	,992	1,008

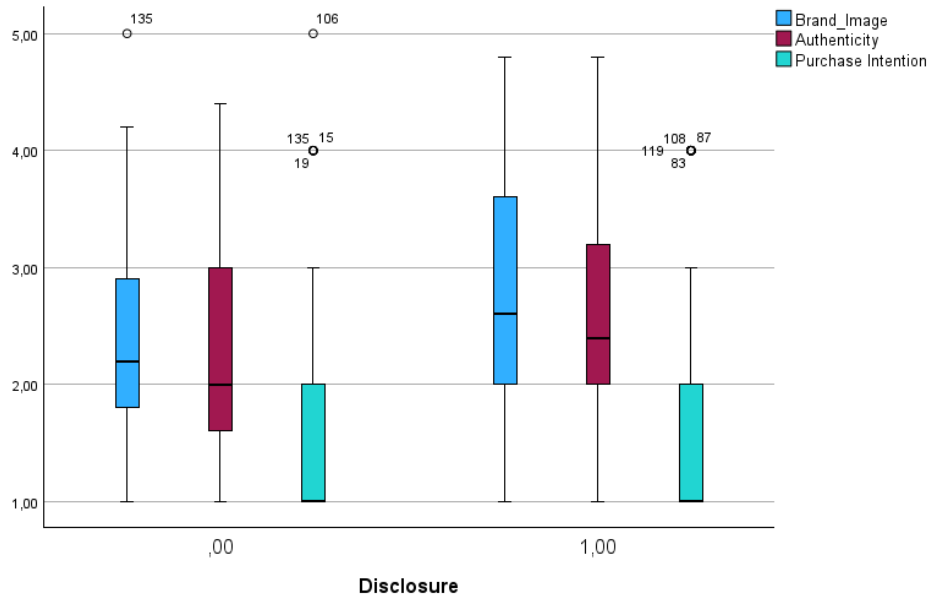
a. Abhängige Variable: Purchase Intention

Koeffizienten^a

Modell		Nicht standardisierte Koeffizienten		Standardisierte Koeffizienten	T	Sig.	Kollinearitätsstatistik	
		Regressionskoeffizient B	Std.-Fehler	Beta			Toleranz	VIF
1	(Konstante)	2,281	,129		17,725	<,001		
	Disclosure	,248	,155	,138	1,604	,111	,992	1,008
	Involvement	,112	,155	,062	,721	,472	,992	1,008

a. Abhängige Variable: Authenticity

Outlier Testing



Independence

Modellzusammenfassung^b

Modell	R	R-Quadrat	Korrigiertes R-Quadrat	Standardfehler des Schätzers	Durbin-Watson-Statistik
1	,156 ^a	,024	,010	,89883	2,174

a. Einflußvariablen : (Konstante), Inv, Disc

b. Abhängige Variable: Auth

Modellzusammenfassung^b

Modell	R	R-Quadrat	Korrigiertes R-Quadrat	Standardfehler des Schätzers	Durbin-Watson-Statistik
1	,227 ^a	,051	,037	,90749	2,220

a. Einflußvariablen : (Konstante), Inv, Disc

b. Abhängige Variable: Br_Img

Modellzusammenfassung^b

Modell	R	R-Quadrat	Korrigiertes R-Quadrat	Standardfehler des Schätzers	Durbin-Watson-Statistik
1	,092 ^a	,008	-,006	1,068	2,315

a. Einflußvariablen : (Konstante), Inv, Disc

b. Abhängige Variable: Purchase Intention

Appendix E: Output for Hypothesis 1a

ANOVA

Authenticity

	Quadratsumme	df	Mittel der Quadrate	F	Sig.
Zwischen den Gruppen	2,265	1	2,265	2,814	,096
Innerhalb der Gruppen	107,870	134	,805		
Gesamt	110,135	135			

Kontrast-Tests

		Kontrast	Kontrastwert	Std.-Fehler	T	df	Sig. (2-seitig)	95% Konfidenzintervall	
								Unterer	Oberer
Authenticity	Gleiche Varianzen voraussetzen	1	,2584	,15402	1,678	134	,096	-,0463	,5630
	Varianzen sind nicht gleich	1	,2584	,15392	1,679	133,255	,096	-,0461	,5628

Kontrast-Tests

		Kontrast	Kontrastwert	Std.-Fehler	T	df	Sig. (2-seitig)	90% Konfidenzintervall	
								Unterer	Oberer
Authenticity	Gleiche Varianzen voraussetzen	1	,2584	,15402	1,678	134	,096	,0033	,5135
	Varianzen sind nicht gleich	1	,2584	,15392	1,679	133,255	,096	,0034	,5133

PROCESS (Authenticity, CI 95 %)

Model: 1

Y: Authenticity

X: Disclosure

W: Involvement

Sample

Size: 136

OUTCOME VARIABLE: Authenticity

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.1836	.0337	.8062	1.5605	3.0000	132.0000	.2020

Model

	coeff	se (HC4)	t	p	LLCI	ULCI
Constant	2,3579	,1453	16,2262	.0000	2,0704	2,6453
Disclosure	.0697	,2206	,3160	,7525	-,3666	,5060
Involvement	-,0549	,2168	-,2531	,8006	-,4836	,3739
Int_1	,3495	,3099	1,1279	,2614	-,2635	,9625

Product terms key:

Int_1: Disclosure x Involvement

Test(s) of highest order unconditional interaction(s):

	R2-chng	F(HC4)	df1	df2	p
X*W	.0093	1,2722	1,0000	132,0000	,2614

Focal predict: Disclosure (X)

Mod var: Involvement (W)

***** BOOTSTRAP RESULTS FOR REGRESSION MODEL PARAMETERS *****

OUTCOME VARIABLE: Authenticity

	Coeff	BootMean	BootSE	BootLLCI	BootULCI
Constant	2.3579	2.3583	.1442	2.0811	2.6488
Disclosure	.0697	.0667	.2193	-.3646	.4941
Involvement	-.0549	-.0511	.2132	-.4567	.3775
Int_1	.3495	.3490	.3083	-.2553	.9479

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
95,0000

Number of bootstrap samples for percentile bootstrap confidence intervals:
5000

PROCESS (Authenticity, CI 90 %)

Model: 1

Y: Authenticity
X: Disclosure
W: Involvement

OUTCOME VARIABLE: Authenticity

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.1836	.0337	.8062	1.5605	3.0000	132.0000	.2020

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
Constant	2.3579	.1453	16.2262	.0000	2.1172	2.5986
Disclosure	.0697	.2206	.3160	.7525	-.2957	.4350
Involvement	-.0549	.2168	-.2531	.8006	-.4139	.3042
Int_1	.3495	.3099	1.1279	.2614	-.1638	.8628

Product terms key:

Int_1: Disclosure x Involvement

Test(s) of highest order unconditional interaction(s):

	R2-chng	F(HC4)	df1	df2	p
X*W	.0093	1.2722	1.0000	132.0000	.2614

Focal predict: Disclosure (X)
Mod var: Involvement (W)

***** BOOTSTRAP RESULTS FOR REGRESSION MODEL PARAMETERS *****

OUTCOME VARIABLE: Authenticity

	Coeff	BootMean	BootSE	BootLLCI	BootULCI
Constant	2.3579	2.3590	.1436	2.1268	2.5944
Disclosure	.0697	.0703	.2211	-.2936	.4379
Involvement	-.0549	-.0572	.2170	-.4135	.3060
Int_1	.3495	.3467	.3081	-.1542	.8535

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
90,0000

Number of bootstrap samples for percentile bootstrap confidence intervals:
5000

Appendix F: Output for Hypothesis 1b

ANOVA

Brand_Image

	Quadratsumme	df	Mittel der Quadrate	F	Sig.
Zwischen den Gruppen	5,072	1	5,072	6,157	,014
Innerhalb der Gruppen	110,394	134	,824		
Gesamt	115,466	135			

Kontrast-Tests

		Kontrast	Kontrastwert	Std.-Fehler	T	df	Sig. (2-seitig)	95% Konfidenzintervall	
								Unterer	Oberer
Brand_Image	Gleiche Varianzen voraussetzen	1	,3866	,15581	2,481	134	,014	,0784	,6948
	Varianzen sind nicht gleich	1	,3866	,15595	2,479	132,433	,014	,0781	,6951

PROCESS (Brand Image, CI 95 %)

Model: 1

Y: Br_Img

X: Disclosure

W: Involvement

Sample

Size: 136

OUTCOME VARIABLE:

Br_Img

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.2279	.0520	.8293	2.5527	3.0000	132.0000	.0583

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
Constant	2.3000	.1366	16.8404	.0000	2.0298	2.5702
Disclosure	.3276	.2217	1.4776	.1419	-.1110	.7661
Involvement	.1182	.2171	.5445	.5870	-.3112	.5476
Int_1	.0876	.3166	.2766	.7825	-.5386	.7138

Product terms key:

Int_1: Disclosure x Involvement

Test(s) of highest order unconditional interaction(s):

	R2-chng	F(HC4)	df1	df2	p
X*W	.0006	.0765	1.0000	132.0000	.7825

Focal predict: Disclosure (X)

Mod var: Involvement (W)

***** BOOTSTRAP RESULTS FOR REGRESSION MODEL PARAMETERS *****

OUTCOME VARIABLE:

Br Img

	Coeff	BootMean	BootSE	BootLLCI	BootULCI	
Constant	2.3000	2.2976	.1372	2.0333	2.5702	
Disclosure	.3276	.3287	.2228	-.1185	.7561	
Involvement	.1182	.1281	.2212	-.3012	.5520	
Int 1	.0876	.0781	.3175	-.5442	.6981	

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
95,0000

Number of bootstrap samples for percentile bootstrap confidence intervals:
5000

PROCESS (Brand Image, CI 90 %)

Model: 1

Y: Br_Img

X: Disclosure

W: Involvement

Sample

Size: 136

OUTCOME VARIABLE: Br_Img

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.2279	.0520	.8293	2.5527	3.0000	132.0000	.0583

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
Constant	2.3000	.1366	16.8404	.0000	2.0738	2.5262
Disclosure	.3276	.2217	1.4776	.1419	-.0397	.6948
Involvement	.1182	.2171	.5445	.5870	-.2414	.4777
Int 1	.0876	.3166	.2766	.7825	-.4368	.6120

Product terms key:

Int_1 : Disclosure x Involvement

Test(s) of highest order unconditional interaction(s):

	R2-chng	F(HC4)	df1	df2	p
X*W	.0006	.0765	1.0000	132.0000	.7825

Focal predict: Disclosure (X)

Mod var: Involvement (W)

***** BOOTSTRAP RESULTS FOR REGRESSION MODEL PARAMETERS *****

OUTCOME VARIABLE: Br Img

	Coeff	BootMean	BootSE	BootLLCI	BootULCI
Constant	2.3000	2.3015	.1373	2.0744	2.5286
Disclosure	.3276	.3278	.2189	-.0369	.6855
Involvement	.1182	.1164	.2178	-.2274	.4803
Int 1	.0876	.0885	.3143	-.4310	.5987

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output: 90,0000
 Number of bootstrap samples for percentile bootstrap confidence intervals: 5000

Appendix G: Output for Hypothesis 1c

ANOVA

Purchase Intention

	Quadratsumme	df	Mittel der Quadrate	F	Sig.
Zwischen den Gruppen	1,217	1	1,217	1,074	,302
Innerhalb der Gruppen	151,812	134	1,133		
Gesamt	153,029	135			

Kontrast-Tests

		Kontrast	Kontrastwert	Std.-Fehler	T	df	Sig. (2-seitig)	95% Konfidenzintervall	
								Unterer	Oberer
Purchase Intention	Gleiche Varianzen voraussetzen	1	,19	,183	1,036	134	,302	-,17	,55
	Varianzen sind nicht gleich	1	,19	,184	1,031	128,563	,304	-,17	,55

PROCESS (Purchase Intention, CI 95 %)

Model: 1

Y: Pur_Int

X: Disclosure

W: Involvement

Sample

Size: 136

OUTCOME VARIABLE:

Pur_Int

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.1004	.0101	1.1476	.5376	3.0000	132.0000	.6573

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
Constant	1.6579	.1423	11.6503	.0000	1.3764	1.9394
Disclosure	.2731	.2444	1.1176	.2658	-.2103	.7566
Involvement	.1300	.2441	.5324	.5953	-.3529	.6129
Int_1	-.1721	.3719	-.4628	.6443	-.9079	.5636

Product terms key:

Int_1: Disclosure x Involvement

Test(s) of highest order unconditional interaction(s):

	R2-chng	F(HC4)	df1	df2	p
X*W	.0016	.2142	1.0000	132.0000	.6443

Focal predict: Disclosure (X)

Mod var: Involvement (W)

***** BOOTSTRAP RESULTS FOR REGRESSION MODEL PARAMETERS *****

OUTCOME VARIABLE: Pur_Int

	coeff	BootMean	BootSE	BootLLCI	BootULCI
Constant	1.6579	1.6611	.1411	1.4000	1.9500
Disclosure	.2731	.2692	.2443	-.2076	.7381
Involvement	.1300	.1339	.2400	-.3214	.6177
Int_1	-.1721	-.1747	.3706	-.8938	.5503

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output: 95,0000

Number of bootstrap samples for percentile bootstrap confidence intervals: 5000

Appendix H: Output for Hypothesis 2a

PROCESS (Brand Image, CI 95 %)

Model: 4

Y: Br_Img

X: Disclosure

M: Authenticity

Sample

Size: 136

OUTCOME VARIABLE: Authenticity

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.1434	.0206	.8050	2.8174	1.0000	134.0000	.0956

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
constant	2.3324	.1072	21.7637	.0000	2.1204	2.5444
Disclosure	.2584	.1539	1.6785	.0956	-.0461	.5628

OUTCOME VARIABLE: Br_Img

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.7985	.6376	.3146	150.3242	2.0000	133.0000	.0000

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
Constant	.4957	.1234	4.0177	.0001	.2516	.7397
Disclosure	.1806	.1004	1.7994	.0742	-.0179	.3792
Authenticity	.7972	.0502	15.8894	.0000	.6979	.8964

Test(s) of X by M interaction:

F(HC4)	df1	df2	p
.4743	1.0000	132.0000	.4922

***** TOTAL EFFECT MODEL *****

OUTCOME VARIABLE: Br_Img

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.2096	.0439	.8238	6.1451	1.0000	134.0000	.0144

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
constant	2.3549	.1067	22.0786	.0000	2.1440	2.5659

Discl	.3866	.1560	2.4789	.0144	.0782	.6951
-------	-------	-------	--------	-------	-------	-------

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y

Effect	se(HC4)	t	p	LLCI	ULCI
.3866	.1560	2.4789	.0144	.0782	.6951

Direct effect of X on Y

Effect	se(HC4)	t	p	LLCI	ULCI
.1806	.1004	1.7994	.0742	-.0179	.3792

Indirect effect(s) of X on Y:

	Effect	BootSE	BootLLCI	BootULCI
Authenticity	.2060	.1232	-.0311	.4555

***** BOOTSTRAP RESULTS FOR REGRESSION MODEL PARAMETERS *****

OUTCOME VARIABLE: Auth

	Coeff	BootMean	BootSE	BootLLCI	BootULCI
constant	2.3324	2.3313	.1069	2.1217	2.5425
Disclosure	.2584	.2621	.1536	-.0389	.5646

OUTCOME VARIABLE: Br_Img

	Coeff	BootMean	BootSE	BootLLCI	BootULCI
constant	.4957	.4962	.1240	.2508	.7446
Disclosure	.1806	.1801	.0996	-.0167	.3752
Authenticity	.7972	.7968	.0508	.6937	.8980

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
95,0000

Number of bootstrap samples for percentile bootstrap confidence intervals:
5000

PROCESS (Brand Image, CI 90 %)

Model: 4

Y: Br_Img

X: Disclosure

M: Authenticity

Sample

Size: 136

OUTCOME VARIABLE: Authenticity

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.1434	.0206	.8050	2.8174	1.0000	134.0000	.0956

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
constant	2.3324	.1072	21.7637	.0000	2.1549	2.5099
Disclosure	.2584	.1539	1.6785	.0956	.0034	.5133

OUTCOME VARIABLE: Br_Img

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.7985	.6376	.3146	150.3242	2.0000	133.0000	.0000

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
constant	.4957	.1234	4.0177	.0001	.2913	.7000
Disclosure	.1806	.1004	1.7994	.0742	.0144	.3469
Authenticity	.7972	.0502	15.8894	.0000	.7141	.8802

Test(s) of X by M interaction:

F(HC4)	df1	df2	p
.4743	1.0000	132.0000	.4922

***** TOTAL EFFECT MODEL *****

OUTCOME VARIABLE: Br_Img

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.2096	.0439	.8238	6.1451	1.0000	134.0000	.0144

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
constant	2.3549	.1067	22.0786	.0000	2.1783	2.5316
Disclosure	.3866	.1560	2.4789	.0144	.1283	.6449

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y

Effect	se(HC4)	t	p	LLCI	ULCI
.3866	.1560	2.4789	.0144	.1283	.6449

Direct effect of X on Y

Effect	se(HC4)	t	p	LLCI	ULCI
.1806	.1004	1.7994	.0742	.0144	.3469

Indirect effect(s) of X on Y:

	Effect	BootSE	BootLLCI	BootULCI
Authenticity	.2060	.1229	.0111	.4127

***** BOOTSTRAP RESULTS FOR REGRESSION MODEL PARAMETERS *****

OUTCOME VARIABLE: Auth

	Coeff	BootMean	BootSE	BootLLCI	BootULCI
Constant	2.3324	2.3338	.1073	2.1570	2.5139
Disclosure	.2584	.2578	.1532	.0143	.5149

OUTCOME VARIABLE: Br_Img

	Coeff	BootMean	BootSE	BootLLCI	BootULCI
constant	.4957	.4962	.1242	.2915	.6981
Disclosure	.1806	.1793	.1008	.0153	.3470
Authenticity	.7972	.7968	.0502	.7141	.8793

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:

90,0000

Number of bootstrap samples for percentile bootstrap confidence intervals:

5000

Appendix I: Output for Hypothesis 2b

PROCESS (Purchase Intention, CI 95 %)

Model: 4

Y: Pur_Int

X: Disclosure

M: Authenticity

Sample

Size: 136

OUTCOME VARIABLE: Authenticity

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.1434	.0206	.8050	2.8174	1.0000	134.0000	.0956

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
constant	2.3324	.1072	21.7637	.0000	2.1204	2.5444
Disclosure	.2584	.1539	1.6785	.0956	-.0461	.5628

OUTCOME VARIABLE: Pur_Int

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	p
.7447	.5546	.5125	71.4409	2.0000	133.0000	.0000

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
constant	-.3356	.1627	-2.0626	.0411	-.6575	-.0138
Disclosure	-.0381	.1289	-.2959	.7677	-.2931	.2168
Authenticity	.8806	.0770	11.4422	.0000	.7284	1.0328

Test(s) of X by M interaction:

F(HC4)	df1	df2	p
.3450	1.0000	132.0000	.5580

***** TOTAL EFFECT MODEL *****

OUTCOME VARIABLE: Pur_Int

Model Summary

R	R-sq	MSE	F(HC4)	df1	df2	P
.8092	.0080	1.1329	1.0629	1.0000	134.0000	.3044

Model

	coeff	se(HC4)	t	p	LLCI	ULCI
constant	1.7183	.1190	14.4453	.0000	1.4830	1.9536
Disclosure	.1894	.1837	1.0310	.3044	-.1739	.5527

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y

Effect	se(HC4)	t	p	LLCI	ULCI
.1894	.1837	1.0310	.3044	-.1739	.5527

Direct effect of X on Y

Effect	se(HC4)	t	p	LLCI	ULCI
-.0381	.1289	-.2959	.7677	-.2931	.2168

Indirect effect(s) of X on Y:

	Effect	BootSE	BootLLCI	BootULCI
--	--------	--------	----------	----------

Authenticity	.2275	.1367	-.0324	.5020
--------------	-------	-------	--------	-------

***** BOOTSTRAP RESULTS FOR REGRESSION MODEL PARAMETERS *****

OUTCOME VARIABLE: Authenticity

	Coeff	BootMean	BootSE	BootLLCI	BootULCI
constant	2.3324	2.3332	.1066	2.1267	2.5494
Disclosure	.2584	.2578	.1533	-.0377	.5659

OUTCOME VARIABLE: Pur_Int

	Coeff	BootMean	BootSE	BootLLCI	BootULCI
constant	-,3356	-,3336	,1609	-,6526	-,0229
Disclosure	-,0381	-,0367	,1288	-,2782	,2277
Authenticity	,8806	,8786	,0762	,7254	1,0236

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:

95,0000

Number of bootstrap samples for percentile bootstrap confidence intervals:

5000