



Ca' Foscari
University
of Venice

**Master's Degree programme
in Innovation and Marketing**

Final Thesis

**Cognitive and Affective Consequences of
Mandatory AI Disclosure in Luxury Advertising**

How Perceived Effort and Moral Disgust Produce Asymmetric
Effects on Brand Authenticity and Willingness to Pay

Supervisor

Ch. Prof. Andreas Hinterhuber

Graduand

Lorenzo Paro

Matriculation Number 889505

Academic Year

2025 / 2026

ABSTRACT

The EU AI Act's mandatory transparency obligations require brands deploying Generative Artificial Intelligence (GenAI) in advertising to disclose AI involvement explicitly, creating a tension with the luxury sector's dependence on human craftsmanship, exclusivity, and brand authenticity. Prior research has examined AI disclosure effects under voluntary conditions in mainstream and hospitality contexts; the consequences of mandatory disclosure within luxury advertising remain empirically unexamined. The present study addresses this gap through a Stimulus-Organism-Response (S-O-R) framework, proposing that mandatory AI disclosure activates two parallel internal mechanisms, a cognitive evaluation of reduced Perceived Effort and an affective response of Moral Disgust, which produce asymmetric consequences for Brand Authenticity and Willingness to Pay (WTP). Consumers' Need for Uniqueness is theorized as a boundary condition moderating the disclosure effect on both mediators. A single-factor between-subjects experiment (N = 145) exposed Italian consumers to a photorealistic advertisement for the fictitious luxury brand MERIDIAN, randomly assigned to an AI-generated or human-photography disclosure condition and analyzed using PLS-SEM. Results confirm that mandatory AI disclosure significantly reduces Brand Authenticity through Perceived Effort as a complementary mediator, whereas the economic penalty on WTP operates exclusively through an indirect-only mediating route. Moral Disgust produces a direct effect on Brand Authenticity but does not extend to WTP, establishing an asymmetric dual-pathway structure. Need for Uniqueness moderation is not supported across any of the four hypothesized pathways. These findings extend dual-pathway consumer response theory to mandatory regulatory disclosure contexts, with direct implications for luxury brand managers navigating GenAI adoption under the EU AI Act.

Keywords: Generative artificial intelligence, Mandatory AI disclosure, Luxury advertising, Brand authenticity, Perceived effort, Moral disgust

RINGRAZIAMENTI

Desidero rivolgere un ringraziamento al Prof. Hinterhuber per la supervisione e le indicazioni che hanno guidato questo lavoro. Un sentito riconoscimento va inoltre a tutte le persone che hanno preso parte al questionario: senza il loro contributo questa analisi non sarebbe esistita. Non mi sarei mai aspettato una partecipazione così ampia, e l'essere riuscito a raggiungere oltre duecento persone soltanto tramite messaggio e passaparola, vedendo quante di loro abbiano dedicato il proprio tempo a compilarlo, mi ha commosso.

Un pensiero più personale va alla mia famiglia, a mia nonna, ai miei zii e ai miei cugini, che mi hanno supportato e sopportato in questi quasi vent'anni di studi. Mi sono sempre impegnato al massimo per portare a casa un ottimo risultato, senza mai esserne obbligato, ma anche per dimostrare che il loro sostegno non è mai stato vano.

Desidero inoltre ringraziare tutti i miei amici, con i quali ho passato bellissimi momenti in questi anni. Devo ammettere che il solo trovarci settimanalmente o sentirci ogni giorno ha aiutato molto. Distrarsi, svagarsi, divertirsi, scherzare e pensare a tutt'altro fuorché allo studio per quelle poche ore è stato, e continuerà a essere, molto importante per me. È probabilmente una delle poche cose a cui non riuscirei a rinunciare facilmente.

Un pensiero riconoscente è rivolto anche ai colleghi universitari, e in particolare ad Anna, con la quale ci siamo aiutati a vicenda specialmente nei momenti più critici e stressanti.

Infine, lascio per ultimo un ringraziamento speciale. Papà, non ci sei più da molto oramai, non hai avuto occasione di festeggiare la prima laurea e purtroppo nemmeno la seconda, ma sono sicuro che ti avrebbe fatto tanto, tanto piacere vedermi arrivare fin qui, e che saresti stato orgoglioso e fiero. Per favore, continua a guardarmi mentre proseguo lungo la mia strada e ricorda: una persona scompare solo quando viene dimenticata da tutti.

Grazie,
Lorenzo

TABLE OF CONTENTS

1. Introduction.....	1
2. Literature Review	6
2.1 Methodology for the Literature Review	6
2.2 The Luxury Sector	16
2.2.1 Characteristics, Values, and Market Positioning of Luxury Brands	16
2.2.2 Human Craftsmanship as a Core Value of Luxury Brands	20
2.2.3 Brand Authenticity and Willingness to Pay in the Luxury Sector	23
2.2.4 Visual and Sensory Dimensions of Luxury Brand Communication	25
2.3 Generative Artificial Intelligence in Marketing.....	28
2.3.1 The Adoption of GenAI in Marketing and Advertising Practice	28
2.3.2 The AI-Made vs. Human-Made Paradox.....	32
2.3.3 The AI Disclosure Effect and the EU AI Act.....	37
2.4 Consumer Psychological Responses to AI Authorship Attribution.....	42
2.4.1 Cognitive Responses to AI Authorship: Effort Heuristic and Perceived Effort	42
2.4.2 Affective Responses to AI Authorship: Moral Disgust and the Uncanny Valley Effect	46
2.5 Individual Differences in Consumer Responses to AI-Generated Content	50
2.5.1 Need for Uniqueness: Theoretical Foundations and Consumer Behavior Relevance	50
3. Hypotheses Development and Key Research Questions	55
3.1 The S-O-R Model and its Adoption in GenAI Luxury Marketing.....	61
3.2 Hypotheses Development	63
3.2.1 AI Disclosure Effects on Willingness to Pay and Brand Authenticity	64

3.2.2 The Cognitive Pathway: Perceived Effort (Mediator 1)	66
3.2.3 The Emotional Pathway: Moral Disgust (Mediator 2)	69
3.2.4 The Moderating Role of the Need for Uniqueness	72
3.3 Summary of the Research Hypotheses	75
4. Methodology	77
4.1 Research Approach and Experimental Design.....	77
4.1.1 Quantitative Research Approach and Experimental Design.....	77
4.1.2 Declaration of AI Tools Utilization	79
4.2 The Stimuli Development and the Image Generation Process	80
4.2.1 The Generative AI Model (FLUX.2 [dev]).....	80
4.2.2 The Image Creation Process	83
4.3 Questionnaire Structure and Measurement Scales.....	88
4.3.1 Questionnaire Structure	88
4.3.2 Measurement Scales	91
4.3.3 Attention and Manipulation Checks	96
4.4 Sample and Data Collection.....	98
4.4.1 Target Population and Sampling Strategy.....	98
4.4.2 Data Preparation and Exclusion Criteria	99
4.4.3 Sample Characteristics.....	102
4.5 Data Analysis Strategy.....	104
5. Data Analysis and Results	109
5.1 Introduction to the PLS-SEM Analysis	109
5.2 Assessment of the Measurement Model (Outer Model).....	112
5.2.1 Indicator Reliability (Outer Loadings).....	113
5.2.2 Internal consistency (Cronbach's Alpha, Composite Reliability)	116

5.2.3 Convergent Validity (Average Variance Extracted).....	118
5.2.4 Discriminant Validity (Cross-loadings, Fornell-Larcker, HTMT)	119
5.2.5 Summary of the Measurement Model Assessment.....	125
5.3 Assessment of the Structural Model (Inner Model)	126
5.3.1 Collinearity Assessment (Variance Inflation Factor).....	129
5.3.2 Path Coefficients and Direct Effects.....	130
5.3.3 Model's Explanatory Power and Effect Size (R^2 , f^2)	138
5.3.4 Parallel Mediation Analysis: Perceived Effort and Moral Disgust.....	142
5.3.5 Moderated Mediation Analysis: Need for Uniqueness	148
5.3.6 Assessment of Control Variables: Luxury Interest and Visual Appeal.....	151
5.4 Analysis of the Willingness to Pay distribution.....	154
5.5 Summary of Hypothesis Testing Results.....	157
6. General Discussion	165
6.1 Overview of Findings	165
6.2 Theoretical Implications.....	168
6.3 Managerial Implications	175
6.4 Contributions to Theory	182
7. Limitations and Future Research	185
8. Conclusions	190
Bibliography.....	194
Appendix.....	208

LIST OF FIGURES

Figure 2.1 Literature Search and Sample Selection Process.....	8
Figure 2.2 Thematic convergence of the literature review domains	54
Figure 3.1 The research model.....	60
Figure 4.1 ComfyUI workflow used to create the GenAI image.....	82
Figure 4.2 A detailed view of the ComfyUI workflow (1)	85
Figure 4.3 A detailed view of the ComfyUI workflow (2)	85
Figure 4.4 Examples of rejected images with AI Artifacts	86
Figure 4.5 The final, unaltered image, selected for this study	87
Figure 4.6a (left) The stimulus depicting the Human Label	88
Figure 4.6b (right) The stimulus depicting the GenAI Label	88
Figure 5.1 The Path Model created in SmartPLS 4	110
Figure 5.2 The Path Model created in SmartPLS 4 with the outer loadings.....	128
Figure 5.3 WTP distribution for N = 145.....	155
Figure 5.4 Final graphical representation of the model	163

LIST OF TABLES

Table 2.1 Summary of the Most Relevant Scientific Literature on GenAI in Marketing and Advertising.....	9
Table 3.1 The most relevant scientific articles for the present study and their relationship to the key research constructs	57
Table 3.2 Overview of the research hypotheses.....	76
Table 4.1 Measurement Scales	92
Table 4.2 Summary of Attention and Manipulation Checks	97
Table 4.3 Summary of Data Cleaning and Exclusion Process.....	101
Table 4.4 Sample Demographic Characteristics (N = 145)	102
Table 4.5 Overview of Constructs, Measurement Items, and Variable Coding for the PLS-SEM Analysis	106
Table 5.1 Outer loadings of the indicators.....	115
Table 5.2 Cronbach's alpha, Composite Reliability, and AVE of the Measurement Model	117
Table 5.3 Cross-loadings of the Measurement Model	120
Table 5.4 Fornell-Larcker of the Measurement Model	122
Table 5.5 HTMT of the Measurement Model	123
Table 5.6 BCa bootstrap confidence intervals of the Measurement Model	125
Table 5.7 Inner VIF Values of the Structural Model	129
Table 5.8 Direct Path Coefficients (β) of the Structural Model.....	131
Table 5.9 Bootstrapping Total Effects results	133
Table 5.10 BCa Bootstrapped Confidence Intervals for the Total Effects of the Structural Model.....	134

Table 5.11 R^2 and R^2_{adj} for the Structural Model.....	139
Table 5.12 f^2 for the Structural Model.....	140
Table 5.13 Specific Indirect Effects of the Structural Model.....	144
Table 5.14 BCa Bootstrapped Confidence Intervals for the Specific Indirect Effects of the Structural Model.....	144
Table 5.15 Total Indirect Effect for the Structural Model.....	147
Table 5.16 Path coefficients for the Moderated Mediation.....	149
Table 5.17 Index of Moderated Mediation.....	150
Table 5.18 Path coefficients and significance of the control variables.....	152
Table 5.19 WTP comparative summary across the Human and the AI-Generated conditions.....	155
Table 5.20 Direction, β , t-statistic, p-values, and hypothesis decision.....	158
Table 5.21 f^2 effect sizes for supported hypotheses	162
Appendix A.....	208
Appendix B.....	210

1. Introduction

The global luxury goods sector has long sustained its commercial architecture on a value proposition that industrial-scale automation cannot replicate: human craftsmanship, brand authenticity, and the deliberate restriction of supply as a mechanism of perceived exclusivity. Estimated at approximately €1.2 trillion by 2018 and characterized by structural resilience, the luxury market's strength derives from creative and artisanal foundations that prior waves of digital technology had largely left intact (Diallo et al., 2021; Xu & Mehta, 2022). The rapid integration of Generative Artificial Intelligence (GenAI) into marketing and advertising production has introduced pressures that prior digital technologies did not. Unlike conventional AI applications such as recommendation engines and predictive analytics, which classify and optimize existing data, GenAI generates novel content across text, images, audio, and video, enabling it to substitute for human creative authorship in ways that prior digital technologies could not (Hermann & Puntoni, 2025). The scale of adoption across the marketing sector is substantial, with nearly one-third of major brand communications anticipated to be AI-generated by 2025 (Brüns & Meißner, 2024), and GenAI projected to add up to USD 275 billion to fashion and luxury operating profits within three to five years (To et al., 2025). The production cost logic driving this adoption is clear: text-to-image models can generate a single advertising image at a fraction of one US dollar, compared with over USD 100 for a freelance designer and thousands for a professional photoshoot (Hartmann et al., 2025). Evidence of this adoption is already visible within the luxury sector: Moncler launched an AI-generated campaign at London Fashion Week 2023 through Maison Meta, and Casablanca, a house publicly committed to shooting its campaigns on photographic film, produced its Spring/Summer 2023 collection through AI generation (To et al., 2025). Nevertheless, social media reception to these campaigns was mixed, and industry observers cautioned about a structural incompatibility with luxury brand value (To et al., 2025). However, the efficiency gains GenAI offers conflict directly with the symbolic and emotional foundations of luxury brand value, as automated scalable production is fundamentally

incompatible with the irreplicable human creative effort upon which the luxury value proposition depends (To et al., 2025; Xu & Mehta, 2022).

The question of whether luxury brands should disclose GenAI involvement in their advertising production is no longer a strategic choice, as it has been resolved at the regulatory level. Regulation (EU) 2024/1689 (EU AI Act), adopted in June 2024 and in force from 1 August 2024, constitutes the first comprehensive legal framework for AI regulation across all 27 EU Member States (European Parliament, 2024a; Hickman et al., 2024). Specifically, under Article 50, deployers of AI systems producing synthetic images, audio, video, or text are required to label outputs as AI-generated in a manner clearly perceptible to consumers, with General-Purpose AI obligations enforceable from August 2025 and full compliance extending to all categories from 2 August 2026 (European Parliament, 2024a; Hickman et al., 2024). Non-compliance carries a maximum financial penalty of up to EUR 35 million or 7% of worldwide annual turnover, whichever is higher (Hickman et al., 2024). At the platform level, operationalization of these mandates has already begun. For instance, Meta requires an AI-info disclosure label on artificially generated advertising content across Instagram and Facebook, and its generative tool automatically adds a visible watermark on all outputs (Hartmann et al., 2025; Heimstad et al., 2025). These developments transform the question from whether to disclose to what consequences mandatory disclosure may produce. The luxury sector constitutes the context of highest theoretical vulnerability to these consequences. Its value architecture rests on human authorship signals that AI disclosure may systematically undermine, and the disclosure of their absence carries evaluative and economic consequences that are more severe than in functionally oriented brand contexts (To et al., 2025; Xu & Mehta, 2022).

Prior research has documented the evaluative and behavioral consequences of AI disclosure across communication contexts, establishing that source disclosure reduces advertising attitude, brand authenticity perceptions, and downstream behavioral intentions through algorithmic aversion and persuasion knowledge activation (Lim & Schmälzle, 2024), and that these penalties are more severe for luxury relative to

mainstream brand contexts (To et al., 2025; Xu & Mehta, 2022). Nevertheless, this body of evidence presents three limitations that the present study is designed to address. First, prior AI disclosure research has been conducted exclusively under voluntary conditions, in which brands choose whether to disclose AI involvement. Consumer responses in this setting are shaped not only by the disclosure itself but also by what that choice communicates about the brand's values and intentions. The mandatory regulatory context established by Article 50 produces a structurally distinct consumer response condition, as the disclosure obligation is externally imposed and its occurrence cannot be attributed to any inferred brand-level decision. Second, the cognitive and affective mechanisms through which disclosure produces its consequences have been examined in isolation rather than as parallel pathways within a unified model. Specifically, To et al. (2025) and Heimstad et al. (2025) examine the perceived effort mechanism without incorporating a concurrent affective pathway, whereas Kirk & Givi (2025) establish moral disgust as a response to AI authorship in personal text-based communications without a concurrent cognitive pathway or a luxury advertising stimulus. Whether the two mechanisms operate simultaneously, and whether they may produce equivalent or asymmetric consequences across relational and economic outcomes, has not been empirically tested. Third, the joint effects of mandatory disclosure on both Brand Authenticity and Willingness to Pay within a single structural model have not been examined for the luxury advertising context, leaving open the question of whether the two constructs respond differently to the same disclosure stimulus. Taken together, these limitations define an empirical gap at the intersection of AI disclosure research and luxury brand management that the present study addresses.

To address this gap, the present study adopts the Stimulus-Organism-Response (S-O-R) framework (Mehrabian & Russell, 1974) to examine how mandatory AI disclosure shapes consumer psychological and economic evaluations in the luxury advertising context. Specifically, the mandatory AI disclosure label is operationalized as the Stimulus, with Perceived Effort and Moral Disgust as parallel Organism-stage mediators representing the cognitive and affective pathways, and Brand Authenticity and Willingness to Pay as the

Response-stage outcomes. Need for Uniqueness is further incorporated as a moderating variable, testing whether consumers with a stronger preference for distinctiveness respond more negatively to the AI disclosure label (Loureiro et al., 2024). The present study is designed to address two interrelated questions: through which psychological mechanisms does mandatory AI disclosure shape consumer evaluations in the luxury advertising context, and whether those mechanisms produce equivalent or asymmetric consequences across evaluative and economic outcomes. The framework is tested through a between-subjects online experiment of $N = 145$ Italian consumers employing the fictitious luxury brand MERIDIAN across two conditions differentiated exclusively by the disclosure label. The advertising stimuli were generated using FLUX.2 [dev], an open-weight, state-of-the-art text-to-image model developed by Black Forest Labs and released in November 2025, representing one of the most technically advanced open-source image generation models available at the time of the study (Black Forest Labs, 2025a). The data are analyzed using Partial Least Squares Structural Equation Modelling (PLS-SEM) following the procedures established by Hair et al. (2017).

The empirical results yield partial and asymmetric support for the proposed framework. AI disclosure produces a confirmed negative total effect on Brand Authenticity and a non-significant total effect on Willingness to Pay. The latter does not reflect an absence of economic consequences but rather a suppression pattern, in which a confirmed negative indirect effect through Perceived Effort is partially offset by a positive direct path. The cognitive pathway through Perceived Effort constitutes the dominant transmission mechanism, producing confirmed negative indirect effects on both Brand Authenticity and Willingness to Pay through two distinct mediation patterns. The affective pathway through Moral Disgust exerts a more limited influence, affecting Brand Authenticity but not Willingness to Pay. Need for Uniqueness moderation is absent across all four tested interaction terms, with both interaction coefficients directionally reversed relative to the theorized amplification pattern.

The present study advances the existing literature in three respects. First, the study provides the first experimental test of a parallel cognitive-affective pathway model of mandatory AI disclosure consequences in luxury advertising, establishing that the two

mechanisms produce asymmetric outcomes. This extends prior single-pathway designs that had examined effort suppression and moral disgust independently rather than in simultaneous competition (Kirk & Givi, 2025; To et al., 2025). Second, the study demonstrates that a non-significant total effect on Willingness to Pay does not indicate the absence of economic consequences. Mediation analysis reveals that a genuine negative economic effect of disclosure exists but is masked at the aggregate level by a small positive direct path. Third, the study specifies the contextual features under which Need for Uniqueness amplification of disclosure penalties does not activate, qualifying the theoretical scope of the consumer heterogeneity mechanism established by Granulo et al. (2021) and Loureiro et al. (2024). These contributions carry immediate practical relevance: with the EU AI Act's transparency obligations for AI-generated advertising content having become enforceable, the mechanisms and consequences documented in the present study represent an operational reality for luxury brands active across the European market rather than a prospective regulatory concern (European Parliament, 2024a; Hickman et al., 2024).

The remainder of this thesis is structured as follows. Chapter 2 develops the theoretical foundations through a semi-systematic literature review addressing the luxury sector's value architecture, GenAI adoption in marketing and advertising, consumer cognitive and affective responses to AI authorship attribution, and individual differences in GenAI content evaluation. Chapter 3 formalizes the conceptual model and the research hypotheses. Chapter 4 details the quantitative experimental design, the stimulus development process, and the PLS-SEM analytical strategy. Chapter 5 presents the measurement model assessment and the structural model results. Chapter 6 discusses the theoretical and managerial implications of the findings and their formal contributions to theory. Chapter 7 addresses the study's limitations and directions for future research. Chapter 8 consolidates the study's theoretical and practical contributions in the context of the mandatory disclosure landscape that motivates the research.

2. Literature Review

2.1 Methodology for the Literature Review

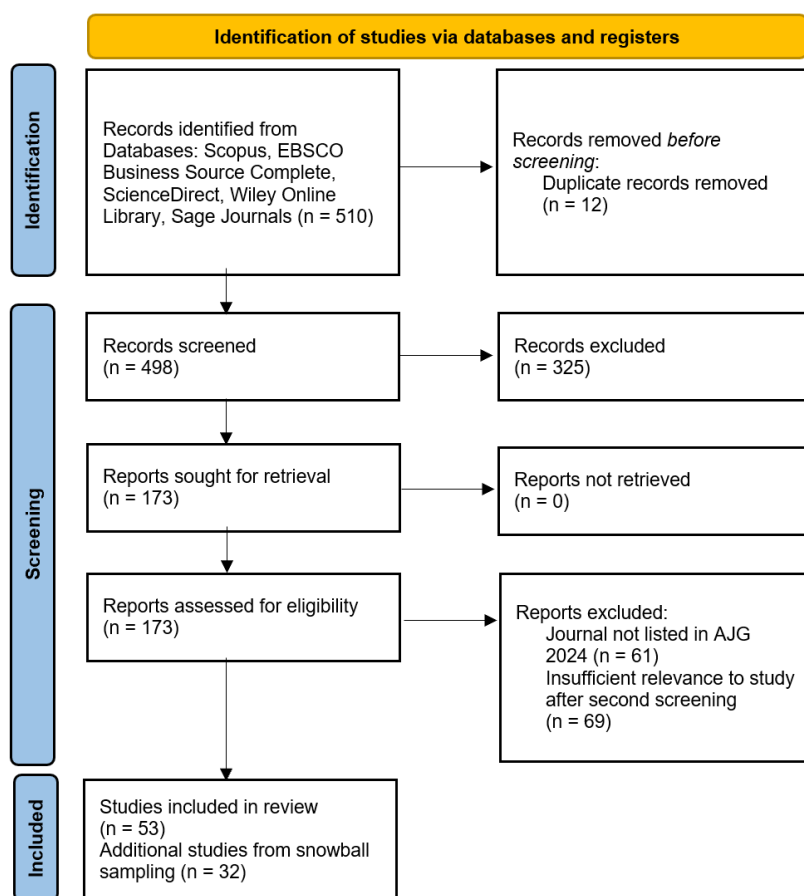
The literature review represents a structured and theoretical examination of the scholarly domains relevant to the central research of the present study, namely Generative AI (GenAI) in marketing and advertising, consumer psychological responses to GenAI content, GenAI advertising within the luxury sector, and individual differences in GenAI consumer behavior. The present literature review adopts a semi-systematic approach, following the typology established by Snyder (2019). This approach is appropriate for the present study as it supports thematic synthesis and theory-grounded integration without requiring the narrow research question characteristic of a systematic review (Snyder, 2019). Accordingly, the main scope of the present literature review is the assessment of the current state of knowledge across each relevant domain to establish the theoretical foundations from which the study's hypotheses are derived in Chapter 3.

The article search was conducted between January and February 2026 across five databases: EBSCO Business Source Complete, ScienceDirect, Wiley Online Library, SpringerLink, and Scopus, selected on the basis of their relevance in the business, marketing, and consumer behavior literature. Five search queries were applied across all databases: ("generative AI" OR "AI-generated content" OR "AI-generated images") AND ("marketing" OR "advertising"); ("AI disclosure") AND ("consumer behavior" OR "brand evaluation" OR "advertising"); ("perceived effort" OR "moral disgust") AND ("artificial intelligence" OR "generative AI"); ("brand authenticity" OR "willingness to pay") AND ("luxury" OR "luxury brands") AND ("artificial intelligence" OR "AI"); and ("need for uniqueness") AND ("AI" OR "artificial intelligence" OR "consumer behavior"). The selected articles were required to be English-language, peer-reviewed journal articles published from 2020 onwards, reflecting the period during which GenAI tools became commercially accessible and the relevant empirical literature emerged. Journal quality was assessed against the Association of Business Schools Academic Journal Guide (AJG) 2024, and only articles indexed within this ranking were considered for inclusion. Peer-

reviewed sources outside the AJG ranking are cited in the broader argument where they provide specific empirical findings not available from ranked equivalents. Purely technical or computer-science journals were excluded. Non-English articles and book chapters were also excluded. Gemini Pro was employed exclusively in the preliminary phase to assist in developing keyword combinations. No AI tool was used to read or summarize any source, with the full-text review conducted by the author, as detailed in Chapter 4. Zotero was used for database deduplication and systematic record management during the search and screening phases. Mendeley was used as the reference manager for in-text citation and bibliography generation throughout the thesis.

The initial database retrieval across the five platforms yielded 510 total records. Following deduplication, 498 records advanced to the screening stage. Title and abstract screening excluded 304 records as non-relevant to the study's domain and a further 21 records on technical or computer-science grounds, leaving 173 records eligible for further screening. Of these, 61 records were excluded for not being indexed in the AJG 2024, reducing the pool to 112 records. A second abstract review excluded a further 25 records, leaving 87 records for a full-text reading. Full-text reading excluded a further 34 records, producing a final database-derived sample of 53 studies. The complete selection process is illustrated in Figure 2.1. The final sample was supplemented by iterative backward snowball sampling of key theoretical and empirical references identified during full-text reading. Several contributions about the study's constructs predate the 2020 inclusion threshold and were therefore not retrievable through the primary database search. Snowball sampling yielded 32 additional sources. In addition, scholarly books were consulted throughout the review. Among these, Simon & Fassnacht's (2019) *Price Management: Strategy, Analysis, Decision, Implementation* represents the primary theoretical reference for the luxury sector. Methodological textbooks on PLS-SEM and construct operationalization were consulted to support the analytical framework and measurement decisions developed in Chapters 4 and 5, with Hair et al.'s (2017) *A Primer on Partial Least Squares Structural Equation Modelling* serving as the principal methodological guide. An overview of the most relevant scientific contributions on GenAI in marketing and advertising is presented in Table 2.1, which details for each study the research context, methodology, key variables, and main findings.

Figure 2.1. Literature Search and Sample Selection Process (Source: PRISMA reporting framework).



The remainder of the chapter is structured as follows. Section 2.2 examines the theoretical foundations of the luxury sector, addressing the defining characteristics and market positioning of luxury goods, the role of human craftsmanship as a core luxury value, brand authenticity and willingness to pay within luxury, and visual and sensory appeal in luxury advertising. Section 2.3 reviews the role of GenAI in marketing and advertising contexts, with particular attention to the AI-made versus human-made paradox and the AI disclosure effect, including its regulatory dimension. Section 2.4 examines consumer psychological responses to AI-generated content, addressing cognitive and emotional responses. Lastly, Section 2.5 addresses the individual differences in consumer responses to AI-generated content, focusing on the Need for Uniqueness construct. Collectively, the theoretical foundations established across these sections provide the evidence from which the conceptual model and the research hypotheses are derived in Chapter 3.

Table 2.1. Summary of the Most Relevant Scientific Literature on GenAI in Marketing and Advertising (Source: the author).

Author(s) & Year	Research Context / Journal	Methodology	Key Variables	Key Findings
To et al. (2025)	AI disclosure effects in luxury vs. mainstream brand advertising <i>Journal of Advertising Research</i>	3 Studies: 1 Field Experiment (Real Ad Campaign) + 2 Lab Experiments	IV: AI Disclosure Med: Perceived Effort, Ad Authenticity Mod: Luxury vs. Mainstream Brand DV: Brand Evaluation	1. Disclosing AI usage leads to negative brand evaluations specifically for luxury brands, while mainstream brands remain largely unaffected. 2. Consumers associate AI with low effort, which directly contradicts the luxury value of craftsmanship, suppressing ad authenticity and brand evaluation.
Jin & Tao (2025)	AI authorship framing and willingness to pay for AI-generated art <i>Psychology and Marketing</i>	4 Lab Experiments (Between-subjects design)	IV: Authorship Framing (Created by AI vs. Generated by AI) Med: Savoring Mod: Visual Style (Abstract vs. Concrete) DV: Willingness to Pay (WTP)	1. The label "Created by AI" triggers a savoring mindset that maintains a high WTP comparable to human-authored art. 2. The label "Generated by AI" implies automated production and significantly reduces WTP.

Author(s) & Year	Research Context / Journal	Methodology	Key Variables	Key Findings
Ali et al. (2025)	GenAI adoption effects on brand perceptions and behavior <i>International Journal of Hospitality Management</i>	Scenario-based experiment (Between-subjects)	IV: GenAI Adoption (AI-generated vs. human-generated content). Med: Brand Authenticity, Brand Image, Self-Brand Congruity. DV: e-WOM, Behavioral Intentions.	1. Perceived brand authenticity, brand image, and self-brand congruity are significantly lower when AI-generated content is disclosed. 2. When AI content is used, self-brand congruity becomes the strongest driver of e-WOM and behavioral intentions, while the traditional effects of brand authenticity diminish.
Heimstad et al. (2025)	Machine heuristics, algorithm aversion, and perceived effort in AI content evaluation <i>Computers in Human Behavior: Artificial Humans</i>	2 Online Experiments (Between-subjects)	IV: Source Label (AI vs. Human) Med: Perceived Effort, Perceived Creativity Mod: Algorithm Aversion DV: Attitude, Quality Judgment	1. Consumers rely on a Machine Heuristic, automatically assuming AI-generated content requires less effort than human-authored content, leading to lower quality and creativity judgments. 2. Consumers with high algorithm aversion penalize AI-attributed content more severely, amplifying the negative effect of source disclosure.

Author(s) & Year	Research Context / Journal	Methodology	Key Variables	Key Findings
Kirk & Givi (2025)	AI authorship, moral disgust, and consumer loyalty in marketing communications <i>Journal of Business Research</i>	7 Preregistered Experiments (Mixed designs)	IV: AI vs Human authorship Med: Moral Disgust, Perceived Authenticity Mod: Reused vs. Original Content DV: Word of Mouth (WOM), Brand Loyalty	1. The AI-Authorship Effect reveals that AI usage in human creative domains evokes moral disgust, leading to a decline in perceived brand authenticity and loyalty. 2. The effect is attenuated when AI is framed as an editor rather than an author, suggesting that perceived creative agency mediates consumer moral responses.
Loureiro et al. (2024)	Consumer self-construal, need for uniqueness, and AI-enabled experience on WTP <i>Psychology and Marketing</i>	Mixed Methods: Qualitative Focus Group + Quantitative Survey (SEM)	IV: Independent Self-Construal Med: AI-enabled Consumer Experience Mod: Need for Uniqueness DV: Avoidance of Similarity, WTP	1. Consumers with an independent self-construal and high need for uniqueness actively avoid AI-enabled products to preserve their distinct identity. 2. This avoidance behavior negatively impacts WTP, particularly among consumers whose self-construal is strongly independent.
Heitmann et al. (2025)	AI advertising performance and brand personality targeting <i>Journal of Marketing</i>	7 Studies: Large-scale Field Tests	IV: Production (AI-tuned vs. Traditional) Mod: Brand Personality	1. AI ads fine-tuned on consumer mindset metrics generally outperform human-produced ads in CTR across a range of product categories.

Author(s) & Year	Research Context / Journal	Methodology	Key Variables	Key Findings
		(Meta/Taboola) + Lab Validation	DV: Customer Engagement, Click Through Rate (CTR)	2. AI is less effective for highly differentiated or personality-specific products, indicating boundary conditions for the performance advantage of AI-generated advertising.
Hartmann et al. (2025)	Visual quality and engagement performance of AI-generated vs. human marketing images <i>International Journal of Research in Marketing</i>	Mixed Methods: Perceptual Image Analysis + Field Experiment (A/B Test)	IV: Image Source (AI-generated vs. Human Photography) DV: Perceived Quality, Realism, Aesthetics, CTR	1. Generative AI produces visuals rated higher in quality, realism, and aesthetics than human photography under blind evaluation conditions. 2. In field tests on social media, AI-generated ads achieved significantly higher engagement than stock photographs.
Park & Ahn (2024)	AI-generated luxury brand personality and consumer purchase intention <i>Journal of Retailing and Consumer Services</i>	2 Lab Experiments (2×2 Between-subjects design)	IV: AI vs Traditional Brand Personality Med: Self-Brand Connection DV: Purchase Intention, Brand Preference	1. AI can successfully replicate the sophisticated visual style of luxury brands, achieving parity with human content in perceived personality fit. 2. Despite visual similarity, consumers exhibit lower purchase intention for AI-generated luxury brand personas, indicating a divergence between aesthetic evaluation and behavioral response.

Author(s) & Year	Research Context / Journal	Methodology	Key Variables	Key Findings
Belanche et al. (2025)	Consumer responses to AI-generated images across utilitarian and hedonic service contexts <i>International Journal of Information Management</i>	3 Experiments (2×2×2 Factorial design)	IV: AI vs Real Images Mod: Service Type (Utilitarian vs Hedonic) DV: Intention to Use, Intention to Recommend	1. Negative consumer reactions to AI-generated images are significantly stronger in hedonic and high-involvement contexts such as tourism and luxury. 2. In utilitarian contexts, AI-generated images are accepted or preferred due to perceived efficiency, indicating a context-dependent boundary condition for AI disclosure effects.
Chung et al. (2025)	Satiation effects of repeated exposure to AI-generated travel imagery <i>Annals of Tourism Research</i>	2 Longitudinal Studies (Repeated exposure)	IV: Image Type (GenAI vs. Real vs. Mixed) DV: Inspiration, Satiation	1. AI-generated images provide high initial inspiration but elicit faster satiation upon repeated exposure relative to real images. 2. Consumer interest in AI-generated content follows an inverted U-shape trajectory, declining rapidly after the initial novelty diminishes.
Hermann & Puntoni (2025)	Conceptual framework for GenAI applications in marketing and principles for ethical design and deployment	Conceptual Paper	Constructs: GenAI Impact Dimensions (Human Enhancement vs. Human Replacement); Marketing Cycle Stage (Research,	1. Proposes a two-dimensional framework distinguishing human-enhancing from human-replacing applications of GenAI across the marketing cycle, demonstrating that the perceived creative agency

Author(s) & Year	Research Context / Journal	Methodology	Key Variables	Key Findings
	<i>Journal of Marketing</i>		Strategy, Marketing Mix); Ethical Principles (ASSURANCE framework: Autonomy, Security, Sustainability, Representativeness, Accountability, Non-biasedness, Crediting, Empowerment)	determines the degree of consumer and societal resistance. 2. Introduces the ASSURANCE principles as a normative framework for ethical GenAI deployment in marketing, with explicit reference to EU AI Act regulatory obligations and mandatory disclosure requirements.
Bui et al. (2024)	Perceived authenticity of AI-generated destination images and its effects on tourist patronage intentions <i>Journal of Destination Marketing & Management</i>	Experimental Research Design (AI-generated vs. human-generated images, labelled vs. non-labelled conditions)	IV: Image Source (AI-generated vs. Human-generated) + Labelling Condition Med: Perceived Authenticity (AI-thenticity), Trust, Tourist-Image Congruence DV: Patronage Intentions	1. Introduces AI-thenticity as a construct capturing the perceived originality and faithful reproduction of AI-generated visual content, demonstrating that perceived authenticity positively influences trust and patronage intentions even for AI-generated stimuli. 2. When visual content is not labelled as AI-generated, AI images are perceived as more authentic than their human-made counterparts. Once disclosed, human-

Author(s) & Year	Research Context / Journal	Methodology	Key Variables	Key Findings
				generated images are perceived as significantly more authentic, establishing disclosure as the critical boundary condition governing authenticity perceptions.
Xu & Mehta (2022)	Consumer responses to AI-designed luxury products across emotional and functional value dimensions <i>Journal of the Academy of Marketing Science</i>	4 Experiments (Between-subjects designs)	IV: Design Source (AI vs. Human Designer) Med: Perceived Brand Essence Mod: Luxury Brand Type (Fashion vs. Automobile); Marketing Message Appeal (Functional vs. Control) DV: Brand Attitude, Purchase Intention	1. AI involvement in luxury product design significantly reduces perceived emotional value and perceived brand essence, leading to lower brand attitude and purchase intention for luxury fashion brands, while leaving non-luxury brand attitudes unaffected. 2. When the functional value of a luxury brand constitutes an equally significant component of its essence, as in luxury automobiles, the negative effect of AI-led design is attenuated and perceived brand essence may increase, establishing luxury value type as the primary boundary condition governing consumer responses to technology adoption in the luxury design process.

Note: IV = Independent Variable; DV = Dependent Variable; Med = Mediator; Mod = Moderator.

2.2 The Luxury Sector

2.2.1 Characteristics, Values, and Market Positioning of Luxury Brands

The definition of luxury has always presented a difficult challenge for researchers and practitioners as the absence of a shared conceptual framework complicates not only the measurement of luxury-related constructs but also the operationalization of brand strategies across different cultural and marketing contexts (Kapferer & Laurent, 2016). Accordingly, rather than converging towards a single classification, the luxury concept operates simultaneously as an economic category, a social signal, and a culturally contingent construct whose boundaries shift across cultural contexts, product categories, and consumer reference groups (Kapferer & Laurent, 2016). In academic and managerial research, several definitional inconsistencies are present: some scholars anchor the concept of luxury to objective attributes such as minimum price thresholds and artisanal production methods, whereas others treat luxury as a subjective, socially constructed perception driven by cultural norms and reference points (Kapferer & Laurent, 2016). Accordingly, a working definition may be derived by examining the attributes that consistently distinguish luxury brands from premium and mass-market equivalents.

Among these attributes, the minimum price threshold represents the most frequently cited objective criterion in luxury literature, as price constitutes an immediately observable and cross-culturally comparable characteristic of the luxury positioning (Kapferer & Laurent, 2016). However, the threshold at which consumers begin to perceive a product as luxury rather than merely premium varies systematically. Kapferer & Laurent (2016) demonstrated extreme dispersion in the point at which consumers perceive a product as luxury, with cross-national variations showing that consumers in Italy and Japan hold significantly higher minimum thresholds (average indices of 151 and 137, respectively) relative to those in the United Kingdom and in the United States (indices of 69 and 65). These thresholds shift across diverse product categories, as consumers apply consistently lower price thresholds to luxury accessories such as pens than to traditional luxury categories such as jewelry (Kapferer & Laurent, 2016). Collectively, these variations confirm that luxury is not a fixed category but a relational one, defined relative to the prevailing price standards of a given product market.

Given these shifting economic boundaries, price positioning remains the most theoretically and operationally distinctive characteristic of luxury goods, separating them from adjacent market segments. Similarly, Simon & Fassnacht (2019) argue that a luxury price position requires sustaining an extremely high price stably and independently of short-term demand fluctuations, as price stability itself functions as a mechanism for protecting brand image and communicating enduring exclusivity. Distinguishing five broad price positions, namely luxury, premium, medium, low, and ultra-low, Simon & Fassnacht (2019) emphasize that luxury and premium differ qualitatively rather than merely in degree. Usually, premium products offer superior quality at elevated but broadly accessible price points, whereas the price in luxury contexts serves primarily as an exclusivity signal (Simon & Fassnacht, 2019). The difference between luxury and premium goods is illustrated by Simon & Fassnacht (2019) in a report of comparative market data: a premium Michael Kors wristwatch retails at approximately \$348, while a luxury A. Lange & Söhne equivalent reaches \$403,000; a standard Hilton hotel room in New York is priced at \$269 compared with the Burj Al Arab Royal Suite in Dubai at \$13,058; and a Ralph Lauren T-shirt at \$79 contrasts with a Prada equivalent at \$740. Consequently, price reductions in luxury markets are strategically destructive, as they damage brand equity by diluting exclusivity perceptions, triggering downward quality inferences, and undermining the symbolic value architecture upon which luxury consumption is premised (Simon & Fassnacht, 2019). This mechanism has been illustrated by the Gucci case, in which the former CEO Patrizio di Marco's attempt to raise handbag prices without an enhancement of brand equity proved commercially damaging, demonstrating that price manipulation alone cannot construct or restore luxury positioning (Simon & Fassnacht, 2019).

Building on the price-demand logic, luxury goods follow an inverted price-demand relationship in which higher prices reinforce rather than suppress desirability, a structural characteristic that distinguishes luxury from all other price positions (Simon & Fassnacht, 2019). This logic has been studied and explained in the literature through three different mechanisms: the Veblen effect, the prestige effect, and the snob effect (Kapferer

& Laurent, 2016; Simon & Fassnacht, 2019). The Veblen effect, originating in Veblen's theory of conspicuous consumption (1899, as cited in Simon & Fassnacht, 2019), holds that demand for luxury goods may increase with price because higher prices provide visible evidence of wealth and social rank, transforming the price of goods into a status signal (Simon & Fassnacht, 2019). The prestige effect operates through a complementary mechanism: in conditions of information asymmetry, where direct quality assessment prior to purchase is difficult or impossible, consumers rely on price as a heuristic for inferring unobservable product quality, with premium pricing justifying assumptions of superior craftsmanship and materials (Simon & Fassnacht, 2019). By contrast, the snob effect captures the decline in desirability that occurs when a luxury product becomes excessively accessible: as a brand's consumer base broadens, exclusive buyers may perceive their symbolic distance from conformist purchasers as diminished, thereby reducing the brand's aspirational value (Amaldoss & Jain, 2005; Kapferer & Laurent, 2016). Collectively, these three effects establish that the price-value relationship in luxury is constitutive of the luxury category itself, with direct implications for how luxury consumers evaluate any perceived threat to a brand's exclusivity or quality signals.

While price represents a core dimension of the luxury market, it constitutes only one component of a broader set of interdependent attributes that collectively distinguish a luxury good from a merely expensive one (Diallo et al., 2021). Controlled production volumes and restricted distribution operate as managed scarcity mechanisms, as illustrated by Ferrari capping its global sales volume to a few thousand units annually, while heritage and brand history serve as long-term authenticity anchors, as observed in luxury groups such as LVMH whose brands average approximately 110 years of history (Simon & Fassnacht, 2019). Superior material quality and technical excellence function as baseline standards, and strong sensory and aesthetic dimensions act as prerequisites for luxury product perception (Simon & Fassnacht, 2019). These attributes are strictly interdependent, with rarity reinforcing prices, heritage reinforcing quality perceptions, and aesthetic appeal reinforcing the symbolic value such that the compromise of any single attribute may risk destabilizing the brand's overall luxury price positioning (Simon & Fassnacht, 2019). This interdependence is theoretically organized by a three-dimensional framework explaining how these attributes generate brand value (Diallo et

al., 2021; Shukla et al., 2015; Simon & Fassnacht, 2019). Within this framework, Shukla et al. (2015) argue that functional value encompasses perceived utility, superior quality, and product desirability; social value captures the brand as a creator of prestige, wealth, and status; while symbolic value reflects self-expression and the desire for uniqueness (Diallo et al., 2021; Simon & Fassnacht, 2019). In luxury consumption, according to Diallo et al. (2021), social and symbolic attributes dominate functional attributes. Conversely, according to Xu & Mehta (2022), this hierarchy does not hold for premium or lower-priced goods, where functional utility tends to prevail and can be seen as the driver that satisfies consumers' needs. Nevertheless, this attribute hierarchy carries direct implications for the present study, as the disclosure of AI-generated advertising operates as the specific mechanism that places the emotional and symbolic value dimensions at risk (To et al., 2025).

Diallo et al. (2021) report that the global luxury market reached an estimated valuation of €1.2 trillion by 2018, characterized by sustained growth and structural resilience relative to mass-market consumer goods, with China's luxury sales alone growing by 20% to €23 billion in that period. Accordingly, market power remains concentrated among a small number of major luxury groups, which largely drive the sector's scale and performance (Simon & Fassnacht, 2019). Consequently, the economic performance of the luxury sector appears grounded in its value architecture, encompassing brand authenticity, exclusivity, human craftsmanship, and heritage (Diallo et al., 2021; Simon & Fassnacht, 2019; Xu & Mehta, 2022). These foundational characteristics imply that any perceived disruption of a core value may trigger commercial and reputational consequences relative to mass-market or premium contexts (Belanche et al., 2025; Kapferer & Laurent, 2016). Beyond these market-level characteristics, individual variation in consumers' general orientation toward the luxury category constitutes a source of heterogeneity in luxury evaluations that warrants acknowledgment. Kapferer & Laurent (2016) demonstrate that luxury sensitivity, defined as the degree of affective engagement with and personal knowledge of luxury goods, varies substantially across consumers and is empirically distinguishable from general product involvement, with

higher sensitivity associated with more favorable baseline evaluations of luxury stimuli, suggesting that individual orientation toward the luxury category may shape how consumers respond to stimuli that threaten the luxury value architecture.

In the present study, the explicit disclosure of AI-generated advertising is examined as a possible disruptor of this value architecture (Kirk & Givi, 2025; To et al., 2025; Xu & Mehta, 2022). Accordingly, the following sections examine the components of the luxury value architecture in greater detail as a theoretical baseline for the GenAI discussion developed in Section 2.3. Specifically, Section 2.2.2 addresses the role of human craftsmanship, Section 2.2.3 examines brand authenticity and willingness to pay, and Section 2.2.4 details the importance of visual and sensory appeal in luxury advertising.

2.2.2 Human Craftsmanship as a Core Value of Luxury Brands

Among the interdependent attributes constituting the luxury value architecture examined in Section 2.2.1, human craftsmanship occupies a theoretically and historically distinctive position (Kapferer & Bastien, 2009). As a value-generating construct, craftsmanship traces its origins back to the pre-industrial artisanal tradition, in which the irreproducibility and uniqueness of handmade goods constituted their primary source of worth, distinguishing them from mass-produced goods (Kapferer & Bastien, 2009; Xu & Mehta, 2022). Accordingly, the literature increasingly treats human craftsmanship not only as an objective production attribute but also as a consumer perception shaped by contextual cues and information signals, such that the value consumers assign to craftsmanship is tied to their beliefs about the production process (Kapferer & Valette-Florence, 2021; To et al., 2025; Wang, 2022). This dimension is particularly relevant for the present study, which operates at the level of consumer responses to disclosed AI-generated advertising production rather than at the level of actual manufacturing methods.

Beyond its association with technical quality and functional excellence, craftsmanship encompasses a narrative of human authorship and creativity that contributes to the symbolic and emotional value dimensions of luxury goods and brands, as identified in Section 2.2.1 (Fuchs et al., 2015; Pedro et al., 2024; Xu & Mehta, 2022). Accordingly, the

presence of a human creator gives luxury goods a sense of authenticity and personal connection, qualities that are identified as irreplaceable elements in the construction of a luxury brand identity (Belanche et al., 2025; Park et al., 2024). Consequently, human craftsmanship functions as a primary antecedent of perceived brand authenticity in luxury contexts, as human effort signals genuineness and originality to consumers evaluating the brand (De Kerviler et al., 2022; Pedro et al., 2024; To et al., 2025). This relationship extends to advertising communications as well as to the product itself: consumers evaluate the authenticity of the advertising as a signal of the brand's broader commitment to its values, such that perceived effort in advertisement production mediates the relationship between production origin disclosure and advertisement level of authenticity perceptions, which in turn affects brand evaluation (To et al., 2025). Consequently, this dual relationship explains why the disclosure of AI-generated advertising constitutes a specific threat to authenticity perceptions, as examined in detail in Section 2.2.3.

To infer product qualities and authenticity, consumers actively utilize perceived labour, time, and skill as a quality heuristic, inferring higher quality and authenticity when significant human effort is perceived, largely independently of any objective quality difference (De Kerviler et al., 2022; Fuchs et al., 2015; Kruger et al., 2004). This mechanism, termed the effort heuristic, operates not only through the direct evaluation of a product's physical attributes but also through mediated signals such as advertising communications, through which consumers infer the degree of dedication invested in the underlying creative process (Kruger et al., 2004). Consequently, perceived effort acts as a central psychological mediator between production-origin cues and downstream consumer evaluations, directly shaping quality judgments, authenticity perceptions, and willingness to pay (De Kerviler et al., 2022; Fuchs et al., 2015; Xu & Mehta, 2022). Conversely, when AI is identified as the source of production, consumers automatically apply a "machine heuristic," assuming that the generative process required minimal time, persistence, or skill (Heimstad et al., 2025; Jung et al., 2025; To et al., 2025). This automated cognitive shortcut systematically suppresses perceived effort, triggering negative consequences for both perceived authenticity and brand evaluation (Belanche et al., 2025; Kirk & Givi, 2025; Lee & Kim, 2024). The suppression of perceived effort persists

even when AI-generated image advertising achieves a quality similar to human photography, demonstrating that the penalty associated with AI disclosure is attributable to the perceived absence of human creativity rather than to any objective lack in visual quality (Hartmann et al., 2025; To et al., 2025). Furthermore, the evaluative consequences of suppressed perceived effort extend beyond cognitive judgments, as the perceived absence of human creative agency may also activate distinctive emotional responses in consumers (Kirk & Givi, 2025; Russell & Giner-Sorolla, 2011). Specifically, this affective pathway operates through moral disgust, a construct examined in detail alongside the broader cognitive and emotional response mechanisms in Section 2.4 (Kirk & Givi, 2025; To et al., 2025).

Beyond its psychological dimension, human craftsmanship also operates as a mechanism that guarantees quality excellence and communicates exclusivity (Simon & Fassnacht, 2019). Luxury brands employ strictly controlled production volumes with vertical supply chain integration to ensure that the artisanal scarcity of products is maintained, reinforcing its premium price positioning (Kapferer & Bastien, 2009; Simon & Fassnacht, 2019; Wang, 2022). While AI-driven production may be highly acceptable or even preferred in functionally-driven utilitarian contexts, where efficiency is prioritized, it conflicts directly with luxury's foundational value proposition of irreplaceable human skill and creativity, as established in Section 2.2.1 (Belanche et al., 2025; Xu & Mehta, 2022). Consequently, this reliance on human craftsmanship establishes strict limitations for technological integration within the luxury sector, as technology and automated production appear to be fundamentally incompatible with the hedonic and symbolic dimensions of luxury consumption (Xu & Mehta, 2022). Accordingly, the literature suggests that craftsmanship constitutes a primary antecedent of perceived brand authenticity and willingness to pay in luxury contexts, as examined in Section 2.2.3 (De Kerviler et al., 2022; To et al., 2025).

The structural incompatibility between artisanal exclusivity and technological automation gives rise to a luxury paradox. This paradox is conceptualized as a conflict

between GenAI capabilities, characterized by seamless automation, infinite scalability, and the elimination of human labour, and luxury's established value architecture founded on human effort, irreplicability, uniqueness, and authentic craftsmanship (To et al., 2025; Xu & Mehta, 2022). Furthermore, this paradox represents a novel challenge specific to GenAI rather than to digital technologies broadly, as prior waves of digitization already augmented the operational efficiency of luxury brands without eliminating human creative agency (Hermann & Puntoni, 2025; Jung et al., 2025). Text-to-image GenAI models possess the capacity to emulate and substitute for human creative authorship in advertising production, posing a threat to the symbolic and emotional value architecture upon which luxury consumption is premised (Hartmann et al., 2025; To et al., 2025). Consequently, consumers may perceive this paradox as a violation of an implicit contract, whose theoretical foundations are examined in more detail in Section 2.3.

2.2.3 Brand Authenticity and Willingness to Pay in the Luxury Sector

Brand authenticity is defined as the subjective genuineness ascribed to a brand, capturing the extent to which it is perceived as authentic, trustworthy, and credible (Morhart et al., 2015). Accordingly, authenticity is understood as a combination of objective brand cues and subjective consumer perceptions, and its assessment is heavily dependent on how individuals interpret brand signals rather than solely on visible facts (Fritz et al., 2017; Morhart et al., 2015). The perception of authenticity also depends on consumers' beliefs about whether brand communications are understood as the product of genuine human emotional states and creative agency (Fuchs et al., 2015; Kirk & Givi, 2025). Furthermore, the luxury context tends to amplify the sensitivity of authenticity perceptions to these cues (To et al., 2025). In luxury contexts, consumers expect true, human-crafted experiences at a level not expected in mainstream brands (De Kerviler et al., 2022; To et al., 2025). Moreover, brand heritage constitutes one of the primary authenticity cues in luxury contexts, providing historical continuity and narrative credibility that tie consumer perceptions to brand genuineness over time, sustained by consistent brand values, traditions, and practices across generations (Dion & Borraz, 2015; Morhart et al., 2015; Xu & Mehta, 2022). Consequently, the literature acknowledges that authenticity is

operationalized across multiple distinct dimensions such as continuity, credibility, integrity, and symbolism, rather than being a unidimensional construct (Fritz et al., 2017; Morhart et al., 2015). Accordingly, luxury brands tend to activate several of these dimensions simultaneously through heritage communication and craftsmanship signals, as established in Section 2.2.2.

Brand authenticity functions as a foundational dimension of luxury brand equity, upon which perceived quality, symbolic value, and emotional resonance heavily depend (Diallo et al., 2021; Morhart et al., 2015). Morhart et al. (2015) argue that authentic brands tend to generate stronger consumer-brand relationships with heightened trust, superior quality attributions, and positive advocacy (word-of-mouth). Furthermore, the consumer trust derived from brand authenticity functions as a mediating construct in high-involvement hedonic purchases, where perceived risk is higher and therefore trust reduces the cognitive burden of committing to a high-price purchase (To et al., 2025). Accordingly, authenticity violations, such as through the disclosure of AI-generated advertising in processes, produce severe consequences in luxury relative to mass-market contexts, as they undermine the core integrity of the brand (Longoni & Cian, 2022; To et al., 2025). When such violations occur, they tend to directly suppress consumers' trust, signaling a breach in the implicit contract between a luxury brand and a consumer (Belanche et al., 2025; Kirk & Givi, 2025). Consequently, these authenticity violations produce a cascading effect in perceived quality and emotional value by eliminating signals of exclusivity and reducing consumers' willingness to pay in luxury contexts (Longoni & Cian, 2022; To et al., 2025).

Willingness to pay is theoretically distinct from purchase intention, as it captures the maximum price a consumer is prepared to pay for a product, reflecting total perceived value rather than an attitudinal disposition to purchase (Schmidt et al., 2024). Furthermore, perceived authenticity operates as an antecedent of willingness to pay in the luxury sector, and this mediation is stronger in luxury than in other consumption contexts because of the dominance of non-functional attributes in luxury value perception, as established in Section 2.2.1 (Diallo et al., 2021; Kapferer & Valette-Florence,

2021). Accordingly, this authenticity-WTP relationship appears particularly pronounced in hedonic and high-involvement consumption contexts, where the dominance of non-functional value dimensions makes consumers' price commitments sensitive to any perceived violation of brand integrity (Diallo et al., 2021; Xu & Mehta, 2022).

Within this theoretical framework, the distinction between hedonic and utilitarian consumption appears relevant for understanding the structural vulnerabilities of luxury brands. Luxury purchasing is characterized by its experiential, symbolic, and affective dimensions, and, consequently, the value proposition rests on non-functional attributes rather than on purely utilitarian benefits (Holbrook & Hirschman, 1982). Since luxury consumption relies heavily on the processing of authenticity cues and emotional resonance, the brand equity of luxury firms tends to rest on the preservation of these intangible assets (Morhart et al., 2015). Consequently, brand authenticity and consumers' willingness to pay in luxury are particularly exposed to any disruption that compromises the perceived integrity of the brand's symbolic offering (Longoni & Cian, 2022; To et al., 2025). A further dimension that warrants examination concerns how luxury brands communicate visually, as the sensory and aesthetic qualities of brand communications affect consumer authenticity evaluations. This relationship is examined in Section 2.2.4.

2.2.4 Visual and Sensory Dimensions of Luxury Brand Communication

Building on the brand authenticity and willingness to pay framework established in Section 2.2.3, a further dimension of the luxury value architecture warrants theoretical examination: the visual and sensory appeal through which luxury brand communications convey, reinforce, and sustain the authenticity signals upon which their brand equity depends. Among the defining attributes of luxury goods, the six-facet framework described by Kapferer (2016), building on Dubois et al. (2001, as cited in Kapferer, 2016), establishes that aesthetics and polysensuality occupy a foundational position of luxury brand communication, grounding visual and sensory appeal as constitutive rather than supplementary dimensions of luxury brands communication.

Accordingly, luxury advertising tends to rely on visual imagery advertising deployed across selective premium media, with visual brand elements such as logos, symbols, and packaging functioning as strategic signals of quality and prestige to both existing and prospective consumers (Keller, 2009; Kim et al., 2016). Furthermore, the concept of art infusion developed by Hagtvedt & Patrick (2008) suggests that aesthetically sophisticated visual content produces a spillover effect of luxury perception independent of its content. This effect suggests that it is the aesthetic communication itself, rather than what the communication depicts or claims, that triggers quality and exclusivity inferences in consumers (Hagtvedt & Patrick, 2008). Consequently, this visual mechanism, operating through a peripheral heuristic, connects directly to the effort heuristic established in Section 2.2.2, whereby consumers infer production investment and brand commitment from visible aesthetic quality. Moreover, luxury brand engagement appears to extend beyond the visual domain and encompass multisensory cues that reinforce the emotional and symbolic value of the consumption experience (Kim et al., 2016). The aesthetic standard of luxury brand communication therefore operates as a threshold condition, such that falling below the expected visual and sensory standard risks undermining the symbolic value architecture that sustains premium pricing and consumer trust (Keller, 2009; Kirk & Givi, 2025; Morhart et al., 2015).

Nevertheless, the literature suggests that the relationship between objective visual quality and perceived brand authenticity may operate as separate constructs rather than as a unified evaluative response to brand communications. Hagtvedt & Patrick (2008) establish that luxury perceptions, triggered by aesthetic cues, depend on the inferred production origin, human skill, and creativity underlying the communication rather than on its objective visual quality alone. This dynamic implies that stimuli of equivalent aesthetic quality may generate divergent authenticity evaluations depending on the perceived production process (Hagtvedt & Patrick, 2008). Recent empirical findings from the GenAI literature provide a direct example of this dynamic in advertising contexts. Under blind evaluation conditions, Hartmann et al. (2025) demonstrate that AI-generated images can surpass human photography in perceived quality, realism, and aesthetic appeal, establishing that objective visual quality is no longer exclusive to human production. Conversely, the explicit disclosure of the artificial origin tends to reverse this

aesthetic preference, with authentic images receiving substantially higher evaluations once the production source is revealed, establishing disclosure rather than visual quality as the critical boundary condition in consumer responses (Califano & Spence, 2024; Chen et al., 2024). Bui et al. (2024) provide direct experimental confirmation of this dynamic through a labelled versus non-labelled between-subjects design in a tourism destination context, demonstrating that AI-generated images are perceived as more authentic than human-generated equivalents when no label is applied, but that once disclosed as AI-generated, human-produced images are perceived as significantly more authentic, confirming that disclosure rather than objective visual quality governs the direction of authenticity evaluations. In luxury contexts specifically, Park & Ahn (2024) demonstrate that AI-generated content achieves aesthetic parity with human brand communications but elicits lower purchase intention among luxury consumers. In synthesis, visual appeal and brand authenticity function as separate constructs in luxury contexts. Objective visual quality constitutes a necessary but insufficient condition for positive consumer evaluations, as authenticity perceptions mediated by production-origin beliefs tend to override purely aesthetic judgments (Califano & Spence, 2024; Hartmann et al., 2025; Park & Ahn, 2024).

The four dimensions developed across Section 2.2, namely premium pricing logic and market positioning (Section 2.2.1), human craftsmanship and the effort heuristic (Section 2.2.2), brand authenticity and willingness to pay (Section 2.2.3), and visual and sensory appeal (Section 2.2.4), collectively constitute an interdependent luxury value architecture in which rarity reinforces price positioning, heritage reinforces quality perceptions, aesthetic quality reinforces authenticity evaluations, and craftsmanship signals reinforce the symbolic and emotional value of the brand as a whole. This interdependence makes the luxury architecture sensitive to disruptions that simultaneously suppress perceived effort, undermine brand authenticity, and separate the effect of favorable visual quality from consumer evaluations. Accordingly, Section 2.3 examines the GenAI landscape as the specific technological development whose adoption and required disclosure interact directly with the established luxury framework.

2.3 Generative Artificial Intelligence in Marketing

2.3.1 The Adoption of GenAI in Marketing and Advertising Practice

Generative Artificial Intelligence (GenAI) constitutes a qualitatively distinct concept within the Artificial Intelligence (AI) landscape, as its capacity to generate novel content, including text, images, audio, and video, differs substantially from the classification and optimization of existing characteristics of conventional AI applications such as recommendation engines and predictive analytics (Hermann & Puntoni, 2025; Kumar et al., 2025). Furthermore, prior to the emergence of GenAI, AI systems operated primarily across mechanical tasks for repetitive automation and thinking tasks for analytical prediction, with feeling tasks (for emotional communication) representing the most advanced and least operationalized domain (Grewal et al., 2025; Huang & Rust, 2021). Accordingly, GenAI is the first category of AI capable of operating simultaneously across all three domains, producing outputs that did not previously exist, including content that appears indistinguishable from human equivalents (Grewal et al., 2025; Hermann & Puntoni, 2025; Puntoni et al., 2021). This distinction introduces human creative authorship as the specific variable being substituted rather than augmented, as GenAI's capacity to emulate and substitute human creative production distinguishes it from all prior waves of marketing automation which augmented operational efficiency without displacing human creative agency (Hermann & Puntoni, 2025).

The shift from a more traditional and back-end AI to GenAI occurred with the global release of ChatGPT in November 2022, followed by models such as DALL-E 3, Midjourney, and Stable Diffusion, which allowed consumers to generate creative outputs without being specialized practitioners (Brüns & Meißner, 2024; Grewal et al., 2025; Kumar et al., 2025). Prasanna & Kushwaha's (2025b) synthesis of 104 peer-reviewed studies confirms that the field has matured from definitional debate to empirical consumer behavior inquiry, with 2024 emerging as the peak publication year, and psychology and cognitive theories, social communication frameworks, and innovation adoption models operating as the dominant theoretical lenses (Prasanna & Kushwaha, 2025a, 2025b). At the macro level, GenAI was projected to contribute USD 463 billion of economic value in 2023 to the marketing sector alone (Chui et al., 2023; Hartmann et al., 2025), with industry forecasts

anticipating that nearly one-third of major brand communications would be generated through AI by 2025 (Ali et al., 2025; Brüns & Meißner, 2024). Furthermore, the economic impact of GenAI across business sectors is projected to exceed USD 1.3 trillion by 2032 (Grewal et al., 2025). For the fashion and luxury sectors specifically, GenAI is forecast to add up to USD 275 billion to operating profits within three to five years (To et al., 2025), a projection that underscores the particular commercial pressure these sectors face to adopt the technology and to comply with the mandatory disclosure obligations it triggers.

To understand the breadth of this adoption across marketing functions, the two-dimensional framework proposed by Hermann & Puntoni (2025) offers a useful organizing framework. This framework categorizes the impact of GenAI along a continuum of human enhancement versus human replacement, operating across marketing research, marketing strategy, and marketing mix, of which implications for creative authorship will be discussed in Section 2.3.2 (Hermann & Puntoni, 2025). Within this landscape, brand adoption of GenAI demonstrates the technology's widespread integration across two clusters. The first cluster represents human replacement and encompasses synthetic fashion photoshoots by Iris von Arnim and Levi Strauss, AI product listings and advertising images by Amazon, and operational cost reduction efforts by Klarna and Cadbury (Brüns & Meißner, 2024; Hermann & Puntoni, 2025; Kumar et al., 2025). These examples suggest that human replacement adoption is not limited to a single industry, but rather represents a broader strategy through which organizations pursue direct cost efficiencies at scale. The second cluster concerns human enhancement applications, in which GenAI augments rather than displaces human creative and communicative labour. For instance, Lexus deployed AI to analyse existing human-authored advertising scripts before generating the screenplay for one of its commercials, while VisitDenmark used it to produce video scripts for destination marketing campaigns (Chen et al., 2024). Grewal et al. (2025) document further instances of measurable performance gains where Vanguard increased LinkedIn conversion rates by 15% through AI-generated advertising copy, Emirates NBD increased credit card leads by 177% through personalized offers, and Unilever reduced customer agent response time by

approximately 90% through AI-generated customer messaging. Collectively, these cases suggest that human enhancement applications tend to amplify existing marketing capabilities without eliminating the human judgment underlying them, and that their value lies not only in the cost reduction per se but also in the measurable performance improvements they generate across diverse functions.

Beyond content generation, the adoption of GenAI extends to personalized consumer communications, representing a strategic adoption application for consumer-brand interaction at the individual level. For personalization, GenAI enables hyper-personalized content and real-time chatbot interactions (Kumar et al., 2025; Mogaji & Jain, 2024). Furthermore, tools such as ChatGPT have demonstrated the capacity to match advertising content to individual consumer profiles across personality traits, exerting a stronger behavioral influence on subsequent consumer decisions than non-personalized equivalents (Brüns & Meißner, 2024). Consequently, GenAI enables the production of multilingual and multimodal content at negligible marginal costs, redistributing core content creation capabilities previously held by specialist agencies to a broader range of market actors, and nudging incumbent firms to redefine their value propositions from execution-based to expertise-driven roles (Wahid et al., 2025). This efficiency driver is evident in production expenses, as GenAI models such as DALL-E 3 can generate a single image at a fraction of one US dollar, versus over USD 100 for a freelance designer and thousands of dollars for professional photoshoots (Hartmann et al., 2025). Although GenAI possesses creative capabilities comparable to humans, the literature suggests that cost reduction, rather than enhancement of creative quality, remains the primary rationale driving its adoption inside organizations (Wahid et al., 2025).

Visual content generation represents the domain in which cost reduction advantages appear most fully realized. Text-to-image models such as DALL-E 3, Midjourney, Stable Diffusion XL, and Adobe Firefly enable the creation of photorealistic advertising imagery, brand photography, and product visualizations directly from natural language descriptions without the need for physical photoshoots (Hartmann et al., 2025; Heitmann et al., 2025). Hartmann et al. (2025), in an evaluation of 10,320 AI-generated images by

254,400 human raters, found that GenAI outperformed human photography on quality, realism, and aesthetics under blind conditions. Furthermore, a replication exercise using Red Bull brand imagery demonstrated that Stable Diffusion XL successfully reproduced the brand's design elements (including the can shape and color palette) without requiring any brand-specific fine-tuning, revealing the technology's capacity for brand-faithful visual production (Hartmann et al., 2025). Accordingly, a field experiment tracking over 173,000 social media impressions yielded AI banner advertisements achieving up to a 50% higher click-through rate than human stock photography, suggesting that AI-generated visual content may function not only as a cost-efficient substitute for human photography but as a potentially superior performer on objective quality and engagement metrics under comparable conditions (Hartmann et al., 2025). Similarly, Califano & Spence (2024) found that GenAI food images were frequently preferred over authentic photographs under blind conditions across unprocessed, processed, and ultra-processed product categories, establishing that the visual quality advantage extends beyond advertising-specific contexts. While Hartmann et al. (2025) establish quality superiority under controlled and impression-based conditions, Heitmann et al. (2025) extend the finding to live campaign performance, demonstrating that AI content fine-tuned on consumer mindset metrics matched or exceeded human-produced content on attention, interest, and engagement in Meta advertisements. Consequently, GenAI democratizes accessibility, as it requires no coding expertise nor large production budgets, enabling firms of all sizes to produce professional advertising imagery at a fraction of prior cost (Cillo & Rubera, 2025; Peres et al., 2023; Wahid et al., 2025). The present study operationalizes the visual advertising context directly through the experimental manipulation of AI-generated advertising stimuli, the methodology for which is detailed in Chapter 4.

Across independent studies, AI-generated visual advertising appears to achieve parity with or superiority over human-produced content on quality evaluations, engagement metrics, and cost efficiency across a broad range of product categories and communication contexts (Hartmann et al., 2025; Heitmann et al., 2025). Nevertheless, the

literature identifies three boundary conditions that qualify the generalizability of performance advantages. First, the effectiveness of AI-generated advertisements is attenuated for highly differentiated or personality-specific brands (Heitmann et al., 2025). Second, Chung et al. (2025) document that AI-generated images and real images follow divergent engagement trajectories under repeated exposure: whereas AI-generated images elicit the highest levels of consumer inspiration at initial exposure before declining steadily along an inverted-U curve, real images reach their peak inspiration at approximately the fourth exposure set before declining along a U-shaped trajectory. This divergence suggests that the quality advantage of AI-generated visuals may be temporally bounded relative to real imagery (Chung et al., 2025). Third, the extensive proliferation of AI-generated content across markets risks eroding distinctiveness and quality differentiation (Wahid et al., 2025). Taken together, these observations suggest that the performance advantages of GenAI in visual advertising do not extend uniformly across all brand and consumer contexts, and that the conditions governing its effectiveness extend beyond visual quality and cost. The theoretical mechanisms underlying these differences are examined in Section 2.3.2.

2.3.2 The AI-Made vs. Human-Made Paradox

Creativity, as argued by Huang & Rust (2021) and Heimstad et al. (2025), operates as part of the human domain and is characterized by intentionality, emotional investment, and original meaning. Accordingly, GenAI is now capable of producing outputs that match or exceed human creative work across written, visual, and advertising domains, while operating through a probabilistic process involving no intentionality or creative struggle (Hartmann et al., 2025; Heimstad et al., 2025; Heitmann et al., 2025). This introduces a structural paradox, as the capability that makes AI-generated outputs indistinguishable from human equivalents simultaneously removes the human origin signal (Hermann & Puntoni, 2025; Kirk & Givi, 2025; Wen & Laporte, 2025). As established in Section 2.2.4, this creates a divergence between objective quality, which may favor AI, and perceived authenticity, which systematically favors human production. Furthermore, GenAI represents the first technological wave capable of operating across feeling tasks,

producing emotional and communicative outputs, rather than solely automating mechanical and analytical processes (Hermann & Puntoni, 2025; Huang & Rust, 2022). Accordingly, this tension appears particularly salient in creative domains, as creativity has historically been among the last human competencies considered immune to automation (Puntoni et al., 2021).

Providing an empirical baseline for this tension, Feng & Song (2026) demonstrate that consumers consistently prefer human-generated artwork even when quality is held constant and the two formats are indistinguishable, suggesting that the devaluation is origin-driven rather than quality-driven. Moreover, Heimstad et al. (2025) and Magni et al. (2024) demonstrate similar patterns in streaming content and visual art respectively, where works labelled as originating from AI received lower effort, creativity, and attitudinal evaluations. Accordingly, this devaluation can be explained by a social identity mechanism rooted in the categorization of human creators as in-group members and AI as outgroup agents (Feng & Song, 2026; Tajfel & Turner, 2004). Therefore, this devaluation reduces empathy and identification, lowering consumers' preference for AI-generated work (Choi & Winterich, 2013; Feng & Song, 2026). While moderating interventions exist to attenuate this categorization, they operate alongside a devaluation mechanism connecting the effort heuristic to the machine heuristic (Heimstad et al., 2025; Kruger et al., 2004). Consumers tend to assign higher quality and value to works perceived as requiring greater human labour through the effort heuristic, and when AI is identified as the creative source the machine heuristic leads them to infer that minimal effort was invested, suppressing quality and creativity evaluations, as explained in detail in Section 2.4.1 (Feng & Song, 2026; Heimstad et al., 2025; Magni et al., 2024). Conversely, this devaluation is sensitive to the perceived nature of the creative process. Jin & Tao (2025) demonstrate that framing the output as "created by AI" maintains willingness to pay comparable to human equivalents, whereas framing it as "generated by AI" significantly reduces willingness to pay. This effect is mediated by a reduction in consumer savoring, defined as the enjoyment of contemplating the product's deeper meaning and perceived creative investment (Jin & Tao, 2025). Consequently, the evaluation of creative works tends to reflect perceived human agency and effort rather than intrinsic properties alone (Hermann & Puntoni, 2025). Furthermore, Kirk & Givi

(2025) document that the use of AI in creative domains may trigger a moral disgust mechanism, as explained in Section 2.4.2.

The divergence between visual equivalence and evaluative authenticity in AI-generated brand communications is captured by the distinction between indexical and iconic authenticity cues established by De Kerviler et al. (2022). Indexical authenticity cues signal that a communication is original and traceable to genuine human creative investment, whereas iconic authenticity cues signal aesthetic resemblance without implying original human creative agency. Within this framework, perceived effort operates as the mediating mechanism, as indexical cues prompt perceptions of greater creative investment, increasing beliefs that the work was produced with genuine care (De Kerviler et al., 2022). Applied to the GenAI context, GenAI may achieve iconic authenticity by visually reproducing a brand's established aesthetic, while systematically failing indexical authenticity which requires traceable human creative origin (De Kerviler et al., 2022). Consequently, GenAI advertising may appear authentic in reproduction terms while remaining inauthentic in origin terms, creating a latent vulnerability that AI disclosure activates. This vulnerability is documented by Brüns & Meißner (2024) who demonstrate that brands adopting GenAI for social media content creation experience reduced perceived brand authenticity among followers, even when the content quality is comparable to human-produced content. This occurs because AI is perceived as lacking the intrinsic human qualities that make brand communications credible and congruent to its identity. Consequently, AI use in brand creative communications suppresses perceived brand authenticity and consumer loyalty by violating the normative expectation of human creative authorship (Brüns & Meißner, 2024; Kirk & Givi, 2025). Furthermore, in fashion design, Lee & Kim (2024) demonstrate that consumers respond more favorably to human-designed clothing than artificially designed equivalents. This happens because AI designs violate consumers' beliefs about authentic creative production, specifically the implicit expectation that creative outputs originate from human intentionality and skill (Lee & Kim, 2024). Nevertheless, Lee & Kim (2024) further document that allowing consumers

to customize AI-designed products partially restores perceived authenticity, suggesting that perceived human agency at any creative stage attenuates the devaluation effect.

Accordingly, To et al. (2025) demonstrate that the disclosure of AI-generated advertising imagery causes negative brand evaluations for luxury brands, while mainstream brands remain largely unaffected, establishing that the devaluation is moderated by the degree to which a brand's value architecture rests on perceived human creative effort as an authenticity signal, as discussed in Section 2.2.2. Conversely, Park & Ahn (2024) demonstrate that AI-generated communications achieve aesthetic parity with luxury brand equivalents but elicit significantly lower purchase intention and reduced self-brand connection among consumers. This divergence connects to the indexical versus iconic authenticity split previously discussed, demonstrating that AI achieves iconic but not indexical authenticity in luxury advertising (De Kerviler et al., 2022; Park & Ahn, 2024). Taken together, these findings suggest that consumer devaluation of AI-generated brand communications operates through two pathways. The first, anticipatory devaluation, arises from consumer inference about what a brand's decision to adopt GenAI signals regarding its commitment to genuine creative investment and operates independently of any explicit disclosure of production origin (Brüns & Meißner, 2024; Lee & Kim, 2024). The second, disclosure-triggered devaluation, is activated by direct knowledge of AI production origin and represents the condition most relevant to the present study, the mechanisms and regulatory context of which are examined in Section 2.3.3.

As established across Sections 2.2.1 through 2.2.4, the luxury sector's value architecture creates a specific vulnerability to GenAI adoption that does not apply equally to more functionally-oriented brands. While functionally-oriented brands may adopt AI without substantially disrupting their core value proposition, luxury brands face a situation in which efficiency gains directly undermine the symbolic and emotional foundations of their premium positioning. Across four experiments employing Louis Vuitton, Burberry, and BMW stimuli, Xu & Mehta (2022) demonstrated that AI involvement in luxury product design significantly reduces perceived emotional value, which is associated with brand essence, heritage, and symbolic resonance, while enhancing functional value or leaving it

intact. For luxury fashion brands, this generates significantly negative consumer attitudes, whereas for luxury automobiles the negative effect tends to be attenuated when the functional value is made salient (Xu & Mehta, 2022). Consequently, the luxury paradox emerges as a structural conflict between GenAI's operational characteristics such as automation, scalability, and elimination of human creative labour, and luxury's value foundations of irreplicable human effort, artisanal exclusivity, and authentic craftsmanship (Jung et al., 2025; To et al., 2025; Xu & Mehta, 2022). This conflict is fundamentally symbolic in nature, as the adoption of AI-generated content communicates to consumers that the brand chose operational efficiency over the human effort and care that they associate with genuine luxury quality and authenticity (Kirk & Givi, 2025; To et al., 2025; Xu & Mehta, 2022). Xu & Mehta (2022) further note that theirs was the first study to examine consumer reactions toward AI-led luxury product design, marking the luxury paradox as a theoretically and managerially distinctive problem without direct precedent in the luxury marketing literature.

To et al. (2025), using real Rolex and Swatch advertising stimuli across three studies, demonstrate that luxury brands disclosing AI-generated imagery experience significantly lower brand evaluations, reduced perceived advertisement authenticity, and negative downstream consumer responses, with perceived effort operating as the mediating pathway. Furthermore, Xu & Mehta (2022) establish that consumers penalize AI in luxury contexts not because they doubt its technical capabilities but because they perceive automated production as inconsistent with the brand's emotional and symbolic essence. Jung et al. (2025) demonstrated that perceived humanness mediates the relationship between AI-generated content and consumer trust, suggesting that restoring perceived humanness may partially attenuate but not eliminate the origin-driven trust gap. Additionally, Belanche et al. (2025) demonstrate that consumer rejection of AI-generated advertising is significantly stronger in hedonic and high-involvement contexts than in utilitarian ones. This is consistent with the emotional versus functional distinction established by Xu & Mehta (2022) and reinforces that the luxury paradox reflects a broader consumer psychology in which the perceived irreplaceability of human creative agency intensifies wherever emotional value is the primary consumption driver (Belanche et al., 2025). Nevertheless, the luxury paradox becomes relevant to consumer

behavior only when AI production origin is made explicitly known, the mechanisms and regulatory context of which are examined in Section 2.3.3.

2.3.3 The AI Disclosure Effect and the EU AI Act

As established in Section 2.3.2, the consumer devaluation of AI-generated brand communications operates through two distinct pathways. The present section examines the second, namely the AI disclosure effect, which is activated by the direct label-based communication of production origin to consumers. The AI disclosure effect is defined as the systematic change in advertising attitude, brand attitude, source credibility, and behavioral intentions arising when AI production origin is explicitly disclosed (Lim & Schmälzle, 2024). The AI disclosure effect is theoretically grounded in two converging bodies of literature. The first, persuasion knowledge theory, suggests that consumers identifying a non-human communication origin activate evaluative coping strategies that largely suppress message effectiveness (Friestad & Wright, 1994; Wortel et al., 2024). The second, source credibility theory, posits that perceived expertise, trustworthiness, and attractiveness directly shape message acceptance (Wortel et al., 2024). Accordingly, AI disclosure appears to activate both theoretical mechanisms simultaneously (Wortel et al., 2024). Nevertheless, Lim & Schmälzle (2024) demonstrate that the direction and magnitude of the effect are moderated by pre-existing attitudes towards AI, establishing individual dispositions as a structural boundary condition. This dynamic introduces a transparency paradox for which consumers appreciate disclosure transparency in principle but exhibit negative attitudinal responses to disclosed content, suggesting that the mandatory disclosure may produce unintended brand equity consequences (Wortel et al., 2024). Accordingly, this tension appears particularly relevant for luxury brands, whose value architecture is most sensitive to the mechanisms that disclosure activates, as established in Section 2.3.2.

The AI disclosure effect operates primarily through two psychological mechanisms, the first of which is algorithmic aversion, representing the tendency to distrust algorithmically generated outputs even when they demonstrably match or outperform human equivalents (Brüns & Meißner, 2024; Dietvorst et al., 2014; Heimstad et al., 2025).

Specifically, Heimstad et al. (2025) demonstrate that in creative contexts, AI-generated content is consistently rated as less valuable, less creative, and less effortful than identical human-labelled content. Brüns & Meißner (2024) extend this to brand communication, showing that algorithmic aversion tends to intensify in identity-driven contexts because automation is perceived as hindering the identity benefits consumers seek from brand engagement. Further evidence establishes algorithmic aversion as a general psychological tendency that disclosure activates. For instance, individuals are less likely to follow medical advice from algorithms rather than doctors, resist AI-driven recruitment processes, and evaluate AI sales personnel more negatively than human equivalents (Brüns & Meißner, 2024; Longoni et al., 2019).

The second psychological mechanism, the machine heuristic, operates as a disclosure-triggered cognitive shortcut, encompassing individuals' pre-existing beliefs about AI capabilities that shape effort and creativity attributions, as established in Section 2.3.2 (Heimstad et al., 2025; Sundar & Kim, 2019; Sundar, 2020). According to Heimstad et al. (2025), AI disclosure triggers a sequential mediation where suppressed perceived effort leads to reduced perceived creativity, which in turn generates negative attitudinal evaluations. Critically, consumers with stronger machine heuristic beliefs attribute more effort and creativity to AI-generated work, attenuating the negative effect and establishing the machine heuristic as both a mechanism and a moderating boundary condition (Heimstad et al., 2025). Furthermore, AI disclosure signals that the communication was produced without authentic human intention, activating persuasion knowledge responses that further suppress advertising and brand attitudes through inferences of manipulative intent (Wortel et al., 2024). Despite the negative effects documented above, content labelled as a human and AI collaboration significantly outperforms AI-only content on attitudinal and behavioral measures as demonstrated by Heimstad et al. (2025), reinforcing the theory of perceived human creative agency as the signal whose absence drives negative disclosure reactions, consistent with the finding in Section 2.3.2 that adding customization cues partially restores perceived human agency (Lee & Kim, 2024).

Empirical evidence across multiple communication contexts consistently documents negative consumer reactions towards AI disclosure on advertising attitude, brand attitude, and downstream behavioral intentions (Lim & Schmälzle, 2024). Wortel et al. (2024) demonstrate that AI disclosure reduces advertising attitude and mediates brand attitude reductions through inferences of manipulative intent, with no direct effect on source credibility. This finding suggests that consumers do not necessarily distrust the brand more upon disclosure but rather find the communication less appealing and potentially manipulative (Wortel et al., 2024). To et al. (2025) establish the luxury-specific severity of the disclosure effect, demonstrating that it is disproportionately more negative for luxury brands relative to mainstream brands, with this asymmetry mediated by perceived effort and advertisement authenticity, as detailed in Sections 2.2.2 and 2.2.3. Similar evidence from the hospitality context is provided by Ali et al. (2025), who demonstrate that AI disclosure significantly reduces perceived brand authenticity, brand image, and self-brand congruity. Self-brand congruity emerges as the strongest driver of behavioral intentions under disclosure, largely displacing the traditional role of brand authenticity, as further detailed in Section 2.5 (Ali et al., 2025).

The disclosure effect is further moderated by individual differences in prior attitudes towards AI, as argued by Lim & Schmälzle (2024), who demonstrate that a negative prior attitude significantly moderates the source disclosure effect. Wortel et al. (2024) corroborate this finding, demonstrating that AI aversion moderates the negative effect of disclosure on brand attitude. Furthermore, no significant differences appear in consumer responses to image-based versus text-based disclosure, establishing that the effects are driven by AI involvement per se (Wortel et al., 2024). The voluntary nature of the disclosure decision can no longer be assumed, as regulatory bodies have mandated the disclosure of AI-generated images, text, audio, and video. Bui et al. (2024) provide converging experimental evidence from a tourism destination context, demonstrating that the disclosure label reverses the direction of authenticity evaluations: AI-generated images are perceived as more authentic than human-generated equivalents when unlabelled, but once the AI origin is disclosed, human-produced images are rated as significantly more authentic, confirming that the label itself governs authenticity responses.

The European Union (EU) Artificial Intelligence (AI) Act, adopted in June 2024 and entering into force on 1 August 2024 as Regulation (EU) 2024/1689, constitutes the first comprehensive legal framework for the regulation of AI systems across all 27 EU Member States, taking a risk-based approach to regulate the lifecycle of different types of AI systems (Hickman et al., 2024). The act employs a risk-based architecture across four tiers encompassing unacceptable risk, high risk, limited risk, and minimal risk, with marketing communications disclosure obligations falling within the limited-risk tier, subject to lighter transparency obligations rather than the more stringent conformity assessments required of high-risk systems (European Parliament, 2024a; Hickman et al., 2024). Non-compliance with the act carries substantial financial consequences, with penalties that may reach up to EUR 35 million or 7 percent of worldwide annual turnover, whichever is higher, establishing the AI Act not as a voluntary framework but as a legally enforceable obligation (Hickman et al., 2024). Consequently, disclosure is no longer a voluntary communicative choice available to the brand but rather a legally imposed condition, such that its effects on consumer perception can be attributed to the informational content of the AI label itself rather than to any inference about the brand's strategic intent. Specifically, Article 50 of the AI Act mandates that deployers of AI systems producing synthetic images, audio, video, or text clearly label outputs as artificially generated in a manner clearly perceptible to consumers, with specific application to visual advertising imagery (European Parliament, 2024a; Wen & Laporte, 2025). This obligation does not apply where the content has undergone human editorial review and a natural or legal person holds editorial responsibility for its publication (European Parliament, 2024a).

GenAI systems such as text-to-image models are classified as limited-risk and subject to transparency obligations under Article 50, while General-Purpose AI model providers carry additional obligations under Article 53, including publishing summaries of copyrighted training data (European Parliament, 2024a). General-Purpose AI provisions became effective from August 2025, with full enforcement commencing from 2 August 2026 (Hickman et al., 2024). At the platform level, Meta requires brands to add an 'AI-info' disclosure label for artificially generated advertising content across Instagram and Facebook, while YouTube has introduced labels for AI-generated or modified realistic content (Heimstad et al., 2025; Wortel et al., 2024). Furthermore, Hartmann et al. (2025)

note that Meta's AI model Imagine automatically embeds a visible watermark, disclosing AI generation on all outputs, and illustrating how platform-level implementation operationalizes disclosure in real advertising practice. The disclosure obligation further reflects a converging global regulatory trajectory, as China's 2023 Provisions on the Administration of Deep Synthesis Internet Information Services introduced mandatory labelling requirements for AI-generated content (Heimstad et al., 2025). Similarly, in the United States, California's Senate Bill 53, signed into law in September 2025 as the Transparency in Frontier AI Act, requires large AI developers to publicly publish a framework documenting their approach to safety and transparency, mandates reporting of critical safety incidents, and establishes whistleblower protections for employees who disclose information about catastrophic risks arising from models' deployment (State of California, 2025). Taken together, these regulatory developments confirm that the EU AI Act represents the most comprehensive and legally binding expression of this global disclosure trend to date (Hickman et al., 2024).

Mandatory AI disclosure transforms the question from whether brands will disclose, to what consequences disclosure produces, and for which brand categories those consequences are most pronounced. As established, the luxury category constitutes the context of highest theoretical and empirical vulnerability to the consequences of mandatory disclosure. The regulatory obligation makes disclosure an unavoidable condition for luxury advertising practice in the EU. Beyond the attitudinal and cognitive responses examined, the disclosure of AI production origin may also trigger affective-moral reactions when consumers perceive the communication as lacking genuine human authorship (Kirk & Givi, 2025). Moreover, the magnitude of these responses is likely to vary across consumers, as individual differences in their need for uniqueness may moderate how strongly origin information shapes evaluative outcomes (Tian et al., 2001). Therefore, the question of primary theoretical interest shifts to the specific mechanisms AI disclosure activates and through which pathways it may produce consequences for luxury brands. Consequently, the cognitive pathway, grounded in the effort heuristic and the machine heuristic, is examined in Section 2.4.1, while the emotional pathway, centered on moral disgust and the uncanny valley effect, is examined in Section 2.4.2. The role of consumer need for uniqueness is addressed in Section 2.5.

2.4 Consumer Psychological Responses to AI Authorship Attribution

2.4.1 Cognitive Responses to AI Authorship: Effort Heuristic and Perceived Effort

Having established in Section 2.3 that AI disclosure operates through algorithmic aversion, the machine heuristic, and persuasion knowledge activation to generate evaluative consequences for luxury brand communications, Section 2.4 turns to examine the consumer-level psychological mechanisms through which those consequences are produced. Among the cognitive responses activated by AI disclosure, perceived effort, defined as the subjective inference consumers form regarding the amount of human labour and creative investment required to produce a given output, occupies a primary position as the initiating cognitive variable through which the evaluative consequences of disclosure are produced at the consumer level (Heimstad et al., 2025; Magni et al., 2024; To et al., 2025). According to Kruger et al. (2004), consumers frequently evaluate creative works and commercial outputs without access to reliable intrinsic quality indicators, relying instead on inferential cues. This dynamic has been termed the effort heuristic and is defined as the tendency for individuals to infer higher quality, value, and liking from works believed to have required greater effort in their production, operating independently of objective differences in the output itself (Kruger et al., 2004). Consequently, across three experiments, Kruger et al. (2004) demonstrate that higher perceived effort consistently generates higher quality evaluations, as established for the luxury context in Section 2.2.2. Furthermore, this heuristic effect appears moderated by ambiguity, operating most strongly when intrinsic quality is difficult to assess (Kruger et al., 2004). This finding suggests a particular theoretical relevance for luxury advertising, where symbolic and experiential dimensions dominate over functional attributes. Extending this dynamic to commercial contexts, Morales (2005) demonstrates across three experiments that consumers tend to reward high-effort firms with greater willingness to pay and more positive overall evaluations at constant quality, with gratitude identified as the mediating mechanism. Specifically, firm effort is interpreted as a controllable and morally significant behavior signaling genuine commitment, which generates positive affect that motivates reward (Morales, 2005). Conversely, when effort is inferred as persuasion-motivated rather than genuinely committed, gratitude is

suppressed and the reward response is eliminated, a vulnerability that holds particular relevance to AI disclosure contexts, as such disclosure labels appear to signal the instrumental production intent that activates this persuasion knowledge reversal (Morales, 2005). Furthermore, Jin & Tao (2025) confirm this dynamic in a visual content context, demonstrating that authorship framing implying automated generation significantly reduces willingness to pay relative to conditions implying human creative investment, even when the visual stimulus is held constant. These foundations define the psychological architecture through which AI source attribution subsequently operates.

Building upon these foundations, Nie & Huang (2023) demonstrate that perceived effort suppression emerges as a generalized cognitive response to automated agents. Across three service experiments, they demonstrate how humanoid service robots are systematically perceived to exert less effort than human employees despite identical output quality. This effort deficit mediates negative service evaluations based on consumer mindset, operating significantly when consumers adopt an abstract mindset but becoming non-significant under a concrete mindset (Nie & Huang, 2023). This devaluation occurs through two distinct mechanisms. First, automated agents are assumed to possess unlimited energy without the cognitive and physical resource constraints that make human effort legible and morally significant, and they produce homogenous standardized outputs rather than individualized responses (Nie & Huang, 2023). Furthermore, Leung et al. (2020) establish effort and talent as independent attribution dimensions to account for why this effort suppression persists even when automated agents perform tasks effectively. Within this framework, effort operates at the process-level, reflecting perceived dedication, whereas talent operates at the outcome-level, reflecting perceived competence (Leung et al., 2020; Nie & Huang, 2023). This distinction suggests that AI-generated content may achieve quality parity with human-produced equivalents, demonstrating talent at the outcome level, while sustaining an effort deficit at the process level, given that producing a high-quality output does not constitute sufficient evidence of effortful human investment (Nie & Huang, 2023). Consequently, this divergence explains how quality parity may exist without evaluative

equivalence, as detailed in Sections 2.2.4 and 2.3.2. Taken together, these sources establish the generalizability of effort suppression as a cognitive response to automated production.

The pattern of effort devaluation documented in automated service contexts extends into AI creative authorship evaluation. Magni et al. (2024), drawing from algorithmic aversion research, demonstrate that individuals consistently perceive GenAI to exert less effort than humans, with this devaluation constituting the primary driver of lower creativity ratings assigned to artificially attributed works. Heimstad et al. (2025) extend this mechanism by providing the formal sequential mediation architecture, building on the machine heuristic framework introduced in Section 2.3.3: AI authorship attribution suppresses perceived effort, which reduces perceived creativity, which in turn shapes overall attitudinal evaluations, with this pathway further extending to behavioral intentions. Furthermore, Heimstad et al. (2025) identify the machine heuristic as an individual-difference moderator within this pathway, such that consumers with stronger pre-existing beliefs about AI capabilities attribute more effort to AI-generated work, attenuating the negative authorship effect and confirming that the effort devaluation response is shaped by individual cognitive frameworks rather than operating uniformly across all consumers. Consequently, despite differences in stimuli and experimental designs, Magni et al. (2024) and Heimstad et al. (2025) establish perceived effort as the central cognitive variable mediating the relationship between AI authorship attribution and downstream evaluations, and confirm it as the mechanism through which AI source attribution produces evaluative consequences.

The downstream evaluative consequences that perceived effort suppression activates operate through two distinct pathways. The first consequence chain operates through brand authenticity, building on the indexical versus iconic distinction introduced in Section 2.3.2 (De Kerviler et al., 2022). Notably, indexical authenticity cues signal original human creative agency, whereas iconic authenticity cues signal aesthetic resemblance without implying original human input (De Kerviler et al., 2022). Perceived effort constitutes the psychological variable through which consumers read indexical

authenticity cues, because a high-effort attribution represents the inference that a human creator invested genuine care and original agency. Accordingly, when AI disclosure suppresses perceived effort, it removes the indexical signal even when visual quality is maintained, such that AI-generated advertising may achieve iconic authenticity while systematically failing indexical authenticity (De Kerviler et al., 2022; To et al., 2025). To et al. (2025) provide the primary empirical confirmation of this effect in the luxury advertising context, demonstrating across multiple studies that AI disclosure reduces perceived effort, which reduces advertisement authenticity, which produces negative advertising evaluations, with this mediated pathway operating significantly in the luxury condition and absent in the mainstream brand condition. The luxury-specific severity of this mediated pathway is connected to the ambiguity moderator established by Kruger et al. (2004). Because luxury quality is difficult to assess through intrinsic examination, as it relies on symbolic and experiential dimensions, the effort heuristic tends to carry greater inferential weight in this category than in mainstream equivalents, making the evaluative penalty activated by AI disclosure disproportionately severe for luxury brands, as established in Sections 2.2.1 and 2.2.2 (Kruger et al., 2004; To et al., 2025).

The second consequence chain operates through the normative-economic pathway identified by Morales (2005), who demonstrated that consumers reward high-effort firms through a gratitude-mediated reciprocity mechanism, and that this mechanism is suppressed when effort is perceived as motivated by persuasion rather than genuinely committed. AI disclosure may activate precisely this suppression condition, as the disclosure label signals automated and instrumental production intent, potentially framing the firm's communicative effort as strategically motivated thereby eliminating the gratitude response that would otherwise support positive evaluations and reward behavior (Morales, 2005). Belanche et al. (2025) provide convergent empirical support, documenting reductions in patronage and purchase-related behavioral intentions under AI disclosure in hedonic and high-involvement contexts, where human creative investment is expected to underpin value perceptions. Accordingly, the evaluative consequences of effort suppression documented across this section appear consequential in the luxury context, as established in Sections 2.2.2 and 2.2.3. Furthermore, the perceived absence of human creative agency also tends to activate a parallel and

theoretically distinct affective response, characterized by moral disgust as a reaction to perceived norm violation in creative authorship, which is examined in Section 2.4.2.

2.4.2 Affective Responses to AI Authorship: Moral Disgust and the Uncanny Valley Effect

Beyond the cognitive responses examined in Section 2.4.1, the disclosure of AI-generated content may also activate a parallel affective dimension that operates independently of deliberate evaluative processes. The consumer behavior literature has consistently documented that affective responses to marketing stimuli can arise with limited cognitive elaboration and exert direct influences on behavioral intentions, producing consequences that are distinct from those attributable to cognitive evaluation alone (Amar et al., 2018; Lerner et al., 2004). Within this affective dimension, moral disgust, which may be broadly understood as a discrete negative emotion elicited by perceived violations of moral norms and principles (Chapman & Anderson, 2013; Kirk & Givi, 2025), emerges as a relevant response in AI authorship contexts, as it tends to motivate avoidance rather than the non-specific evaluative penalties associated with generalized negative affect (Giner-Sorolla & Chapman, 2017; Russell & Giner-Sorolla, 2011).

Moral disgust is elicited by perceived violations of purity, authenticity, or divinity norms and is distinguishable from moral anger in that disgust arises from evidence of bad moral character independently of harmful consequences, while moral anger responds to act wrongness (Giner-Sorolla & Chapman, 2017; Rozin et al., 1999; Russell & Giner-Sorolla, 2011). Accordingly, this distinction appears to be central in the AI authorship context, where the violation concerns the nature of the producer rather than the quality of the output. Furthermore, disgust towards a perceived violation may spread to associated entities, carrying direct implications for brand evaluation (Chapman & Anderson, 2013; Kirk & Givi, 2025). Moral disgust has received growing attention in the marketing literature as a construct relevant to brand inauthenticity and ethical norm violations, and the emergence of AI-generated creative content represents a new and theoretically significant domain of application in which its mechanisms are beginning to be examined systematically.

The AI-authorship effect constitutes the primary empirical phenomenon connecting AI disclosure to moral disgust in creative contexts (Kirk & Givi, 2025). Across seven pre-registered experiments, Kirk & Givi (2025) document that consumers recognizing a creative communication as AI-generated experience moral disgust, producing reductions in word of mouth and customer loyalty. This mechanism operates through the authenticity congruence requirement, which dictates that genuine creative communication requires alignment between a creator's internal emotional state and its external expression (Kirk & Givi, 2025). Because AI systems are perceived as incapable of possessing the internal emotional states that genuine creative communication requires, GenAI content is interpreted not as a lower-quality substitute but as an inauthentic act of expression. Furthermore, Kirk & Givi (2025) specify that this effect tends to be stronger for emotional and hedonic communications than for utilitarian content, and is activated specifically when human emotional authenticity constitutes a normative expectation, as established in Section 2.2.3. Nevertheless, Kirk & Givi (2025) identify two moderation conditions: framing AI as an editor rather than an autonomous author, and positioning AI as an autonomous social entity. The first condition attenuates disgust by preserving the perception of human creative agency as the origin of the communication, while the second condition attenuates disgust by positioning AI as an autonomous communicative entity in its own right so that consumers do not apply the expectation of human authenticity in the first place, removing the normative comparison that generates the violation (Kirk & Givi, 2025).

Consequently, both moderating conditions confirm that the moral disgust response is tied to violated expectations of human authorship rather than to AI involvement per se, such that disgust is activated not by the use of AI as a production tool but by the inference that a communication is presented as carrying human emotional authenticity that an artificial system is fundamentally incapable of possessing (Kirk & Givi, 2025). This mechanism is demonstrated by Wei et al. (2025), who show that moral reactions to AI involvement in creative contexts depend on which psychological process is activated, and when consumers attribute low agency to the AI source, moral condemnation is reduced. This dynamic confirms that the disgust response is driven by violated expectations of human authorship. Accordingly, as established in Section 2.2.2, luxury brand communications are

expected to express human creative commitment and emotional resonance, and it is precisely the perceived absence of these properties in AI-generated content that activates the moral disgust response. Furthermore, Belanche et al. (2025) provide complementary evidence that AI-generated visual content in high-involvement contexts is frequently perceived as impersonal and lacking genuine emotional resonance.

Alongside moral disgust, a complementary and distinct source of affective discomfort in GenAI visual content operates at the perceptual level through the uncanny valley effect. Mori (1970) proposed that as artificial entities grow increasingly human-like in appearance, consumer affective responses initially improve but then reverse sharply into discomfort and rejection once resemblance becomes near-perfect yet detectably imperfect (Kirk & Givi, 2025; Mori et al., 2012). Califano & Spence (2024) document that AI-generated food images perceived as too perfect tend to elicit reduced authenticity and appeal, while Sorosrungruang et al. (2024) demonstrate how virtual influencers resembling near-human entities may generate alienating responses. Furthermore, Hartmann et al. (2025) document a hyperrealism paradox, demonstrating that AI-generated content may achieve higher realism ratings than human-produced equivalents under blind evaluation, yet activate uncanny valley discomfort once AI origin is disclosed. Consequently, while moral disgust is triggered by the knowledge of artificial authorship as a norm violation, the uncanny valley effect operates as a response to visual properties that may arise independently of disclosure. Nevertheless, the two mechanisms are theoretically separable but may operate simultaneously and reinforce one another when AI disclosure coincides with near-realistic visual content. Accordingly, the visual quality threshold established in Section 2.2.4 appears relevant to whether AI-generated stimuli fall above or within the uncanny valley threshold. The uncanny valley effect is therefore acknowledged here as a theoretically separable mechanism that may interact with moral disgust under specific visual conditions, rather than as a distinct and independently testable response.

Having established the mechanisms through which AI disclosure activates moral disgust and, in visually realistic contexts, uncanny valley discomfort, the downstream behavioral consequences that these affective responses produce warrant closer examination, particularly the avoidance and spreading dynamics that distinguish moral disgust from other negative evaluative responses. Amar et al. (2018) develop this argument demonstrating that moral disgust toward perceived counterfeiting spreads via psychological contagion to contaminate genuine products resembling the counterfeit, operating through the Law of Contagion and the Law of Similarity. This contagion logic implies that disgust activated by an inauthentic AI-generated communication may spread to contaminate broader brand authenticity perceptions, a consequence particularly significant for luxury brands, as established in Section 2.2.3. Notably, Kirk & Givi (2025) document the moral licensing reversal, establishing that human authors tend to be penalized more severely than AI authors for identical inauthentic communications. This reversal confirms that moral disgust arises from violated expectations of human authenticity rather than from inauthenticity per se, distinguishing it from a generic anti-AI bias. Kirk & Givi (2025) establish that these behavioral consequences operate through reductions in perceived brand authenticity, such that moral disgust toward an AI-generated communication tends to reduce authenticity perceptions, which subsequently suppress loyalty and positive word of mouth. Consequently, brand authenticity functions as the intermediate mechanism through which the affective response to AI disclosure produces its broader consequences for brand equity (Kirk & Givi, 2025).

The contagion dynamic documented by Amar et al. (2018) clarifies the specific route through which moral disgust produces consequences for brand authenticity perceptions. Because the brand is perceived as the responsible communicative agent, disgust elicited by a norm-violating communication tends to spread to contaminate the consumer's broader evaluation of the brand itself, reducing authenticity perceptions even when the brand's heritage and product quality remain unchanged (Kirk & Givi, 2025). Romani et al. (2012) confirm that disgust-adjacent negative brand emotions produce measurable downstream penalties, including brand switching intentions, mediated by reductions in the evaluative integrity attributed to the brand. Consequently, moral disgust toward AI-generated communication functions as the initiating event in a sequence through which

brand authenticity perceptions are degraded producing measurable behavioral consequences, including brand switching and negative word of mouth (Romani et al., 2012). Further evidence by Narayanan (2025) demonstrates that AI-generated images reduce perceived authenticity via congruity violation, suggesting that the authenticity penalty associated with AI production detection is not restricted to text-based communications. Furthermore, Wortel et al. (2024), as argued in Section 2.3.3, suggest that the moral disgust and persuasion knowledge mechanisms are complementary but distinct and may produce combined penalties when both are activated simultaneously. Nevertheless, the magnitude of the cognitive and affective responses examined in Section 2.4 may vary across consumers with different dispositional orientations toward authenticity and luxury symbolism. This dynamic is examined through the concept of need for uniqueness in Section 2.5.

2.5 Individual Differences in Consumer Responses to AI-Generated Content

2.5.1 Need for Uniqueness: Theoretical Foundations and Consumer Behavior Relevance

While Section 2.4 documents the general psychological processes through which AI disclosure affects consumer evaluations, individual difference variables systematically shape how stimuli are interpreted and evaluated, such that the magnitude of these effects is unlikely to be uniform across all consumers. This heterogeneity motivates the examination of the need for uniqueness as an individual difference variable with direct relevance to AI disclosure contexts. Snyder & Fromkin (1977) introduced the need for uniqueness as a positive striving for differentness relative to others, activated when perceived similarity becomes psychologically aversive. Accordingly, high-need-for-uniqueness individuals tend to exhibit lower conformity responsiveness, greater willingness to risk disapproval, and stronger autonomy drives (Snyder & Fromkin, 1977). Notably, Snyder & Fromkin (1977) demonstrated that uniqueness motivation is both situationally reactive, intensifying when individuals perceive excessive similarity to

others, and dispositionally stable, meaning that people differ reliably in how strongly they feel this drive regardless of context. This dual character predicts not merely whether consumers will respond to standardization cues but how strongly and consistently they will do so across contexts. By operationalizing these foundations in the consumer behavior literature, Tian et al. (2001) establish the construct as the trait pursuit of differentness through the acquisition, utilization, and disposition of consumer goods to develop and enhance personal and social identity. Consequently, to satisfy this need, consumers frequently acquire and display specific products or associate with distinct brands to signal their differentiation and individuality (Tian et al., 2001). Furthermore, this trait-stable but behaviorally context-dependent orientation suggests a specific contextual sensitivity, such that consumers with a high-need-for-uniqueness appear specifically responsive to cues signaling standardization or the erosion of symbolic distinctiveness (Tian et al., 2001). This dynamic allows the construct to be directly applicable to AI disclosure contexts, where disclosure functions precisely as such a standardization cue (Bui et al., 2024; Tian et al., 2001).

Tian et al. (2001) operationalized the need for uniqueness as three empirically distinguishable dimensions: creative choice counterconformity, unpopular choice counterconformity, and avoidance of similarity, each producing different behavioral consequences that justify treating them as independent mechanisms. Specifically, avoidance of similarity is defined as the tendency to lose interest in possessions and brand associations once these become mainstream and cease to provide differentiated value (Tian et al., 2001). Consequently, avoidance of similarity emerges as the most theoretically relevant dimension to standardized content contexts, as it operates as a reactive rather than a proactive mechanism (Tian et al., 2001). Avoidance of similarity is activated when previously distinctive associations become generalized, which may trigger the withdrawal of positive evaluations and reduced willingness to pay (Loureiro et al., 2024; Tian et al., 2001). This dynamic is grounded in Belk's (1988) extended self framework, which establishes that consumers regard their possessions and brand associations as reflections and extensions of their identities, such that any erosion of a brand's distinctiveness constitutes a threat to the self (Tian et al., 2001). Aaker & Schmitt (2001) demonstrate that independent self-construal consumers hold differentiation-

oriented consumption preferences and are disproportionately sensitive to standardization cues while interdependent self-construal consumers prefer assimilation-oriented symbols. Loureiro et al. (2024) confirm this empirically in an AI-enabled consumption context, demonstrating that independent self-construal consumers experience an amplified avoidance response with downstream consequences for willingness to pay when exposed to automated standardization. Nevertheless, it appears that avoidance of similarity, amplified among independent self-construal consumers, represents the dimension most directly activated by the homogenization signals embedded in AI disclosure (Loureiro et al., 2024).

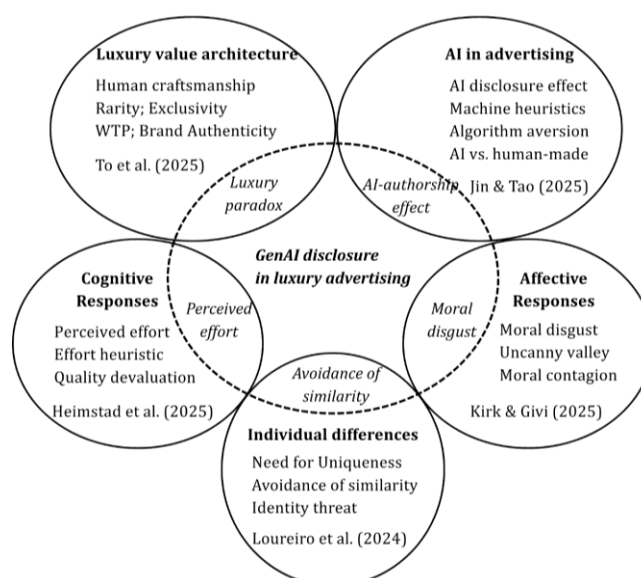
The luxury consumption domain represents the most theoretically significant context for the activation of need for uniqueness because luxury goods and communications are consumed precisely as vehicles for social differentiation and identity expression, making any perceived erosion of exclusivity directly threatening to the consumer's symbolic investment in the brand. For high-need-for-uniqueness consumers, luxury goods function as primary instruments of social differentiation, with rarity, price barriers, and human craftsmanship operating as exclusivity cues that distinguish their consumption choices, as established in Sections 2.2.1 and 2.2.2 (Loureiro et al., 2024; Tian et al., 2001). Within Belk's (1988) extended self framework, luxury brand associations function as components of the consumer's identity, such that perceived mass production threatens their self-expressive function rather than reducing aesthetic satisfaction, a threat particularly severe for high-need-for-uniqueness consumers whose self-concept depends on the distinctiveness signals that luxury brands provide. Furthermore, GenAI enables the scalable production of visually equivalent advertising at near-zero marginal cost, directly undermining the scarcity signal on which luxury brand equity depends (Simon & Fassnacht, 2019). Consequently, AI disclosure signals not only automated production but the erasure of human creative scarcity that anchors the brand's identity-differentiation function. As established in Section 2.2.4, visual quality operates as a necessary but not sufficient condition for positive luxury evaluations, and symbolic differentiation operates independently of visual quality parity. Nevertheless, for high-need-for-uniqueness

consumers, the perception of AI-generated output as standardized directly triggers avoidance of similarity, producing a downstream willingness-to-pay effect that tends to be amplified in the luxury context (Loureiro et al., 2024).

The intersection of need for uniqueness with AI-generated content has begun to receive empirical attention, with evidence that high-need-for-uniqueness consumers are disproportionately sensitive to the standardization and homogenization signals embedded in AI disclosure. Loureiro et al. (2024) and Wahid et al. (2025) argue that GenAI tools deployed at scale rely on shared training data distributions, producing outputs that tend toward common aesthetic conventions that reduce creative distinctiveness. Consequently, this homogenization constitutes a direct trigger for the avoidance of similarity mechanism in high-need-for-uniqueness consumers. Furthermore, Loureiro et al. (2024) demonstrate that high-need-for-uniqueness consumers with independent self-construal respond to artificially enabled experiences through heightened avoidance of similarity with downstream willingness-to-pay consequences. However, their empirical context involves branded AI tools rather than AI-generated advertising stimuli specifically but, nevertheless, the avoidance of similarity mechanism operates through the same perceived homogenization logic and is theoretically applicable to AI disclosure in advertising contexts (Loureiro et al., 2024). Furthermore, according to Granulo et al. (2021), need for uniqueness moderates the interaction between production mode and symbolic motives. Notably, high-need-for-uniqueness evaluative responses to GenAI reverse when content is perceived as undifferentiated, which suggests that the response is driven by the standardization signal rather than AI involvement per se (Loureiro et al., 2024; Sands et al., 2022). Accordingly, the effort devaluation and moral disgust responses established in Sections 2.4.1 and 2.4.2 may both be amplified among high-need-for-uniqueness consumers, as the suppression of perceived human creative effort and the inference of communicative inauthenticity are the cues that signal the erosion of exclusivity and distinctiveness on which luxury brands depend.

The arguments developed across this chapter provide a theoretical account of why the disclosure of AI-generated content in luxury advertising is likely to produce negative consumer evaluations. The luxury sector, as documented in Section 2.2, is defined by a value architecture in which rarity, human craftsmanship, and brand authenticity function not merely as product attributes but as the symbolic foundations of consumer identity investment and willingness to pay premium prices. Sections 2.3 and 2.4 documented the psychological mechanisms through which AI disclosure may disrupt this value architecture, identifying a cognitive pathway in which the perceived absence of human creative effort suppresses brand authenticity perceptions, and an affective pathway in which the inference of communicative inauthenticity generates moral disgust with further consequences for brand equity and consumer loyalty. Section 2.5 introduced the need for uniqueness as a source of systematic consumer heterogeneity, suggesting that these responses may be particularly pronounced among consumers whose identity investment in luxury brands depends on their exclusivity and human creative scarcity. The thematic intersections across these domains, and their collective convergence on the study are represented in Figure 2.2. Taken together, these theoretical arguments provide the basis for a formal model of consumer responses to AI disclosure in luxury advertising, which is formalized in Chapter 3.

Figure 2.2. Thematic convergence of the literature review domains (Source: the author).



3. Hypotheses Development and Key Research Questions

The preceding chapter established the theoretical landscape within which the present study is situated, addressing the luxury sector's value architecture and its foundational dependence on human craftsmanship and brand authenticity, the structural conflict between GenAI adoption and the luxury value proposition, and the cognitive and affective mechanisms through which AI authorship attribution shapes consumer evaluations. As summarized in Table 2.1, prior empirical work has examined GenAI's consequences for brand communications across a range of contexts, yet the specific mechanisms through which mandatory disclosure of AI-generated content affects consumer responses in luxury advertising remain empirically unexamined. Research conducted under voluntary conditions has established that AI source disclosure reduces advertising attitude and brand authenticity perceptions through algorithmic aversion and persuasion knowledge activation (Lim & Schmälzle, 2024), and the cognitive and affective response pathways have been studied predominantly in isolation rather than as parallel mechanisms within a unified structural model (Heimstad et al., 2025; Kirk & Givi, 2025; To et al., 2025).

Consequently, the empirical literature on GenAI's consumer consequences has concentrated predominantly in service and hospitality contexts, while the luxury advertising sector, in which such consequences are arguably most commercially significant, has remained comparatively underexplored. For instance, Ali et al. (2025) found that in the restaurant industry, the disclosure of social media AI-generated content significantly diminishes perceived brand authenticity and image authenticity, forcing consumers to rely more on internal self-brand congruity to form their behavioral intentions. In a similar way, Belanche et al. (2025) demonstrated that consumer rejections of AI-generated promotional images are particularly severe in high-involvement and hedonic services, such as tourism, whereas utilitarian services remain largely unaffected. Furthermore, in the context of creative media, Heimstad et al. (2025) discovered that audiences systematically devalue AI-generated output because they apply a machine heuristic, automatically assuming that AI exerts less effort than a human creator and subsequently damaging perceived quality and attitudes. However, while these studies

clarify the impact of GenAI in service and mainstream contexts, the same cannot be said for tangible products. This gap is particularly relevant for the high-end luxury segment, a market that is intrinsically tied to values such as human craftsmanship, exclusivity, and uniqueness, where the disclosure of AI-generated content may be a disruptor of perceived value (To et al., 2025).

To explicitly position this study within the current academic landscape, Table 3.1 summarizes the most relevant scientific research articles regarding AI disclosure in marketing, consumer behavior, and the luxury sector, and highlights how the present research directly addresses these critical literature gaps. Luxury brands are fundamentally differentiated by an interdependent value architecture encompassing premium price positioning, controlled scarcity, human craftsmanship, heritage, and a distinctive capacity to generate experiential and symbolic value that mass-market equivalents cannot replicate (Kapferer & Bastien, 2009; Simon & Fassnacht, 2019). Within this architecture, the emotional value derived from human creative agency occupies a privileged position: it is the primary driver of perceived brand essence in emotionally-oriented luxury categories such as fashion and accessories, and constitutes the principal justification for the price premium that luxury consumers are willing to allocate (Xu & Mehta, 2022). Authenticity, exclusivity, and human craftsmanship collectively give luxury products and services a unique and personal identity that cannot be standardized or automated without disrupting the symbolic foundations of the brand's value proposition, as detailed in Section 2.2.1 (Simon & Fassnacht, 2019; To et al., 2025).

Table 3.1. The most relevant scientific articles for the present study and their relationship to the key research constructs (Source: the author).

Author(s) & Year	Theoretical Framework	Research Context	AI Disclosure	Luxury Context	Perceived Effort	Moral Disgust	Brand Authenticity	Willingness to Pay	Need for Uniqueness
To et al. (2025)	Authenticity Theory & Effort Heuristic	Disclosing AI-generated images in luxury vs. mainstream advertising	✓	✓	✓		✓		
Kirk & Givi (2025)	Authenticity Theory	Emotional marketing communications by AI vs. Humans	✓			✓	✓		
Loureiro et al. (2024)	Consumer Uniqueness & Self-Construal Theories	AI-enabled consumer experiences and personalized recommendations	✓					✓	✓
Ali et al. (2025)	S-O-R (Stimulus-Organism-Response)	Impact of GenAI social media content creation on consumer behavior in restaurants	✓				✓		

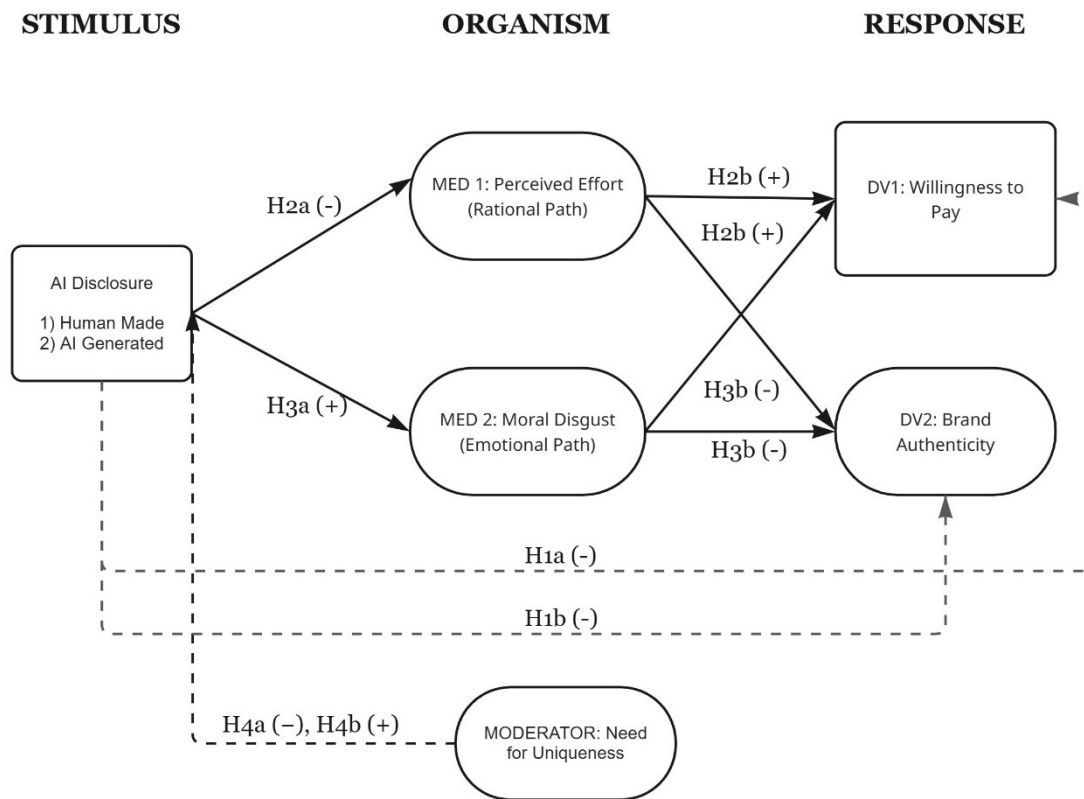
Author(s) & Year	Theoretical Framework	Research Context	AI Disclosure	Luxury Context	Perceived Effort	Moral Disgust	Brand Authenticity	Willingness to Pay	Need for Uniqueness
Jin & Tao (2025)	Savoring Theory & Construal-Level Theory	Framing AI visual products as "Created" vs. "Generated"	✓					✓	
Heimstad et al. (2025)	Algorithm Aversion & Machine Heuristic	Evaluations of creative output (effort and creativity) by AI	✓		✓				
Xu & Mehta (2022)	Authenticity Theory & Psychology of Automation	AI-designed luxury products across emotional and functional value dimensions	✓	✓			✓		
This Study	S-O-R (Stimulus-Organism-Response) Model	How the disclosure of GenAI visual content impacts consumer WTP and Brand Authenticity within the high-end luxury sector	✓	✓	✓	✓	✓	✓	✓

As highlighted in Table 3.1, extant research has primarily examined the psychological and economic consequences of using GenAI advertising in isolated clusters. While some scholars have successfully identified cognitive penalties associated with AI adoption, such as the heuristic assumption of lower perceived effort, other studies have documented the affective and emotional backlash that it can trigger, such as the moral disgust (Heimstad et al., 2025; Kirk & Givi, 2025; To et al., 2025). Consequently, it becomes evident that previous research has failed to integrate these distinct cognitive and emotional reactions to the luxury paradox into a single and comprehensive model. Furthermore, the specific impact of recently mandated AI transparency disclosure on economic and psychological outcomes, such as consumers' willingness to pay and perceived brand authenticity, remains largely unexplored. This limitation is particularly true within sector-specific contexts, such as high-end luxury advertising.

To address the central research questions and bridge the existing literature gap, this study adopts a comprehensive dual-path mediation model grounded in the Stimulus-Organism-Response (S-O-R) framework. This framework has been chosen because it is highly effective for examining how technology-mediated cues influence internal evaluations and subsequent behaviors. For example, Ali et al. (2025) successfully applied this framework in their investigation of GenAI adoption and brand authenticity in the hospitality sector.

The research model, illustrated in Figure 3.1, structurally links the mandated disclosure of AI-generated content (Stimulus) with the consumers' internal psychological processing (Organism). More specifically, the internal psychological process is uniquely divided into two parallel paths: the first connected to a cognitive evaluation of the stimulus (Perceived Effort), and the second connected to an emotional reaction (Moral Disgust). Ultimately, these mechanisms should predict the behavioral and psychological outcomes (Response) of the AI disclosure, namely the consumers' Willingness to Pay (WTP) and perceived Brand Authenticity. Additionally, the model accounts for individual consumer traits by testing how the Need for Uniqueness acts as a moderating variable by influencing the Organism structural relationships as demonstrated by Loureiro et al. (2024).

Figure 3.1. The research model (Source: the author).



Note: The figure illustrates the direct structural paths, mediation pathways, and moderation effects (H1a–H4b) within the S-O-R framework. For visual clarity, hypotheses representing overall or conditional indirect effects (H2, H3, H4c, H4d) are not drawn as separate arrows, as they are mathematically derived from the explicit interactions shown above. Additionally, control variables (Visual Appeal and Luxury Interest) are omitted from this diagram to prevent visual clutter, but they are integrated into the structural model estimated in Chapter 5.

The remainder of the chapter is structured as follows. First, Section 3.1 presents a detailed explanation of the S-O-R framework, justifying its foundational role in mapping consumer reactions to AI-generated advertising. Second, Section 3.2 provides an in-depth breakdown of the specific constructs and formally develops the research hypotheses. Lastly, Section 3.3 provides a summary of the research hypotheses.

3.1 The S-O-R Model and its Adoption in GenAI Luxury Marketing

To systematically investigate the luxury paradox, a robust theoretical framework is required to map the psychological processes through which consumers react to technological cues. Classical consumer behavior models have traditionally relied on a simplistic Stimulus-Response (S-R) framework, in which the consumer mind is treated as a "black box" and purchasing behavior is assumed to follow directly from objective product attributes such as price, performance, and functional features (Mehrabian & Russell, 1974). However, this Stimulus-Response logic fails to capture the complexity of actual consumer decision-making, particularly in high-involvement and hedonic contexts where subjective perceptions, affective reactions, and symbolic interpretations play a key role alongside rational evaluation. In the luxury segment specifically, where consumption is driven primarily by emotional and symbolic value rather than functional utility, a purely Stimulus-Response logic is inadequate, as the same objective stimulus may generate divergent psychological and behavioral outcomes depending on the consumer's prior expectations, identity investment, and contextual evaluative frame (Simon & Fassnacht, 2019; Xu & Mehta, 2022). Therefore, what stands out from the literature is the importance of understanding and exploring these psychological variables. To do so, this study adopts an evolution of the classical S-R model called the Stimulus-Organism-Response (S-O-R) framework, which was originally developed by Mehrabian & Russell (1974). This framework posits that an external environmental stimulus (S) affects an individual's internal cognitive and emotional state (O), which subsequently drives their behavioral and psychological responses (R). The S-O-R framework is particularly suited to understand the dual cognitive and emotional pathways that shape consumers' interactions with new developing technologies (Wu & Li, 2018).

The S-O-R model has recently gained popularity among the GenAI literature, especially in the field of marketing research, due to its efficacy in explaining how technology-mediated communications shape consumer perceptions (Ali et al., 2025; Prasanna & Kushwaha, 2025b). Ali et al. (2025), for instance, demonstrate that when GenAI fully automates rather than merely assists brand content creation, perceived brand authenticity, brand image, and self-brand congruity are all significantly reduced at the Organism stage, with

self-brand congruity emerging as the dominant driver of downstream e-word-of-mouth and behavioral intentions at the Response stage, illustrating how a single technology-mediated stimulus activates a chain of internal evaluations that ultimately governs consumer behavior in the hospitality sector. Moreover, a systematic review of GenAI marketing literature emphasizes that the S-O-R model is highly effective for customized marketing contexts, since it is able to explicitly capture how algorithmic and artificial stimuli trigger complex emotional and cognitive behavioral shifts (Prasanna & Kushwaha, 2025b).

Building on this theoretical framework, the current study rejects the notion of evaluating an AI-generated advertisement by utilizing a utilitarian approach. Specifically, this study adopts the S-O-R model to develop its hypotheses (H1 through H4) by directly mapping the model's components to the specific variables as previously highlighted in Figure 3.1. The framework's components are operationalized as follows.

The Stimulus (S) has been identified as the mandated transparency cue within luxury visual advertisement. More specifically, the disclosure of the label that identifies the visual advertising as "Generated entirely by Artificial Intelligence" or created by "Traditional Photography" functions as the independent exogenous variable. Within the S-O-R framework, this label acts as the external environmental cue that initiates the consumers' evaluation process.

The Organism (O) dimension of the framework represents the internal psychological processing of the consumers. A core strength of the S-O-R framework is its ability to distinguish different mediating mechanisms that operate simultaneously within individuals (Wu & Li, 2018). Therefore, this study proposes a dual-path mediation model where the Organism is divided into a cognitive pathway, represented by the Perceived Effort attributed to the creation of the advertisement, and an affective pathway, represented by the Moral Disgust that arises when AI is utilized in a context traditionally defined by human craftsmanship.

Lastly, the Response (R) stage of the framework seeks to capture the economic and psychological outcomes of consumers driven by the Organism (O) stage. The present study aims to understand how the initial disclosure of an AI label in the luxury advertisement segment impacts consumers' WTP and perceived Brand Authenticity for a luxury brand. Furthermore, a boundary condition has been established to understand whether the effect of the disclosure may be attenuated or amplified. This is operationalized as the consumers' Need for Uniqueness, a variable important for the luxury sector as discussed in Section 2.5.1 (Loureiro et al., 2024). This variable acts as a moderating factor, influencing the structural relationships among the Stimulus, Organism, and Response. By structuring this study through the S-O-R paradigm, this study ensures that the complex mechanism and reactions to the usage of GenAI for luxury advertising are not reduced to a simplistic cause-and-effect model. Instead, this framework provides the grounding basis to define how and why AI disclosure impacts the luxury segment.

3.2 Hypotheses Development

The following sections formalize the theoretical relationships underlying consumers' reactions to GenAI in luxury advertising. Specifically, these sections map the five core constructs of the research model. The latter positions the AI-authorship disclosure label as the external Stimulus, which triggers an internal Organism evaluation characterized by two parallel pathways: a cognitive evaluation driven by Perceived Effort, and an affective reaction caused by Moral Disgust. These internal psychological mechanisms are the drivers of consumer's WTP and perceived Brand Authenticity. Furthermore, consumers' Need for Uniqueness acts as a boundary condition moderating the pathways from the Stimulus to the Organism. To empirically isolate these effects, the experimental stimulus is operationalized using a fictitious luxury brand named "MERIDIAN." This strictly controls for pre-existing brand equity and consumer loyalty biases (the methodological approach of which is explained in Chapter 4).

3.2.1 AI Disclosure Effects on Willingness to Pay and Brand Authenticity

Kirk & Givi (2025) and To et al. (2025) argue that the rapid integration of GenAI in marketing communications has prompted regulatory bodies worldwide to mandate the explicit disclosure of AI-generated content. Examples of such mandates come from the European Parliament through the recent AI Act and, in the United States, from California's Transparency in Frontier AI Act, signed into law in September 2025 (European Parliament, 2024b; Hickman et al., 2024; State of California, 2025). In the context of the present research, the mandated disclosure of AI-generated content (through visual labels) represents the primary and external stimulus necessary to operationalize an S-O-R framework in the field of marketing advertising. To test the proposed framework, the stimulus is operationalized through a fictitious luxury brand named "MERIDIAN." As previously highlighted by Ali et al. (2025) and Kirk & Givi (2025), the methodological rationale for this choice is that utilizing an existing luxury brand would introduce confounding variables, such as pre-existing brand equity and familiarity. The presence of these variables would compromise the causal isolation of the AI disclosure effect on consumers' WTP and Brand Authenticity. The detailed procedure regarding the stimulus design and the specific AI-generated image creation process are established in Chapter 4.

However, the capacity of contemporary GenAI models to produce visually compelling advertising content does not protect luxury brands from the consequences of disclosure. Hartmann et al. (2025) demonstrate that AI-generated images can match or exceed human photography in perceived visual quality and realism under blind evaluation conditions. Nevertheless, Califano & Spence (2024) and Park & Ahn (2024) establish that once AI production origin is made explicit through disclosure, consumer evaluations deteriorate significantly, with authentic human-produced imagery receiving substantially higher evaluations. This pattern confirms that disclosure, rather than objective visual quality, constitutes the critical boundary condition governing consumer responses to AI-generated luxury advertising content. In their studies, To et al. (2025), Ali et al. (2025), and Kirk & Givi (2025) demonstrated how brand authenticity, defined as the subjective genuineness ascribed to a brand, and specifically the extent to which the brand is

perceived as authentic, trustworthy, and credible, is one of the primary constructs systematically penalized by the disclosure of AI-generated content (Morhart et al., 2015). Results of these studies show that when consumers are made aware that a marketing communication is AI-generated, the content and the brand are perceived as less authentic because machines are believed to lack genuine internal emotional states and human intentionality (Kirk & Givi, 2025). To et al. (2025) demonstrated that this penalty effect is stronger in luxury contexts, where consumers explicitly expect true and real brand experiences, crafted by human artisans and with higher effort compared to non-luxury brands. Empirical evidence demonstrated that the disclosure of AI-generated content through labels in luxury advertising drastically reduces perceived advertisement authenticity and brand authenticity, while mainstream brands and advertising remain largely unaffected by the same disclosure mechanism (To et al., 2025). Therefore, this study posits that the mere disclosure of AI-generated visual marketing content through labels systematically deprives luxury advertising content of the expected human factor, damaging the authenticity of the brand.

Alongside the devaluation of brand authenticity, the AI-disclosure mechanism may trigger economic consequences by lowering consumers' WTP, as demonstrated by Belanche et al. (2025). WTP is established in consumer behavior and pricing literature as the maximum financial amount that an individual is willing to allocate for a specific good or service (Hinterhuber & Liozu, 2017; Simon & Fassnacht, 2019). Belanche et al. (2025) show how, in high-involvement and hedonic purchasing decisions (such as the luxury segment), the artificiality signaled by AI disclosure disrupts the cognitive processes of consumers, leading to diminished commercial results. Furthermore, Jin & Tao (2025) confirm this dynamic in a visual content context, demonstrating that authorship framing implying automated generation significantly reduces willingness to pay relative to conditions implying human creative investment, even when the visual stimulus is held constant, establishing that the willingness to pay reduction operates through the perceived absence of human creative effort rather than through any objective difference in visual quality. These findings are consistent with the evidence of Xu & Mehta (2022), who demonstrate across four studies, that AI involvement in luxury product design significantly reduces perceived emotional value and perceived brand essence for emotionally-driven luxury

categories such as fashion, while leaving functional value largely intact. Accordingly, Xu & Mehta (2022) establish that this devaluation is moderated by the degree to which a luxury brand's essence is grounded in emotional rather than functional value, confirming that the penalty is most pronounced in contexts where human creative agency constitutes the primary source of brand worth. Given that luxury advertising communications are evaluated by consumers as signals of the brand's creative investment and emotional authenticity, the same logic extends to the advertising production context examined in the present study.

Therefore, rather than a signal of innovation, the disclosure of AI-generated visual content in luxury advertising acts as a penalty cue for consumers (To et al., 2025). By simultaneously signaling a lack of human effort and creativity and by inducing consumers to think about automated production, beliefs of lower effort arise. These dual psychological (Brand Authenticity) and financial (WTP) effects actively damage luxury brands. Based on this theoretical rationale, the following hypotheses are formally proposed:

H1a: The disclosure of AI-generated visual content negatively impacts consumers' Willingness to Pay compared to traditional human photography;

H1b: The disclosure of AI-generated visual content negatively impacts perceived Brand Authenticity compared to traditional human photography.

It is important to note that H1a and H1b refer to the total effect of GenAI disclosure on the two dependent variables, without specifying at this stage the psychological mechanism through which such an effect operates. The decomposition of this total effect into its cognitive and emotional components is formalized in the mediation hypotheses presented in the subsequent sections (H2 and H3).

3.2.2 The Cognitive Pathway: Perceived Effort (Mediator 1)

The S-O-R framework depicted in Figure 3.1 establishes that within the Organism (O), the external stimulus of an AI disclosure label triggers a cognitive evaluation process. This

cognitive pathway is governed by the consumers' assessment of the labour invested in the advertising creation process, operationalized in this study as Perceived Effort. The theoretical foundation of the Perceived Effort pathway rests on the effort heuristic, a well-established cognitive mechanism whereby consumers assign higher aesthetic, monetary, and quality value to outputs they perceive as having required extensive human time, skill, and creative investment (Kruger et al., 2004). This heuristic operates not only through direct product evaluation but also through mediated signals such as advertising communications, through which consumers infer the degree of dedication invested in the underlying creative process. Perceived effort thereby functions as a psychological signal of commitment and care, shaping downstream quality judgments, authenticity perceptions, and WTP across product and service categories (De Kerviler et al., 2022; Heimstad et al., 2025; Kruger et al., 2004).

To et al. (2025) found that whenever consumers process the explicit disclosure of AI-generated content, they systematically attribute less effort in the creation process. This cognitive devaluation is driven by what the literature terms the machine heuristic, a mental shortcut through which consumers automatically assume that AI-generated outputs require minimal time, creativity, or human persistence, because automated systems are perceived as operating effortlessly and as inherently oriented toward standardization and efficiency rather than genuine creative investment (Heimstad et al., 2025; To et al., 2025). Because this heuristic is activated automatically upon exposure to an AI authorship label, the disclosure itself is sufficient to suppress perceived effort independently of the actual visual quality of the advertising stimulus (Hartmann et al., 2025; Magni et al., 2024; To et al., 2025). The devaluation of AI-generated advertising content is particularly detrimental for the luxury sector, a market in which the premium price point and the prestige associated with the brand are inextricably linked to meticulous craftsmanship and human labour (To et al., 2025). Therefore, in high-end contexts, cues that signal a higher perceived effort are also antecedents of higher perceived brand authenticity, whereas a lack of such effort undermines the brand's genuine expression (De Kerviler et al., 2022). By validating this mechanism in the luxury advertising context, To et al. (2025) demonstrated that the disclosure of AI-generated images directly makes consumers perceive that luxury brands invested less effort into

their marketing campaigns. Subsequently, by violating core luxury expectations of exclusivity and human craftsmanship, negative perceived effort is shown to reduce and damage brand authenticity.

Furthermore, this perception of an effortless process implies an automated mass production approach, diminishing the creative depth and rarity of the product (Wahid et al., 2025). Based on the findings of Kruger et al. (2004), empirical research demonstrates that consumers inherently exhibit a higher WTP when they perceive a high degree of effort was invested in a product's creation. This mechanism is reinforced by Morales (2005), who demonstrates that consumers reward high-effort firms through a gratitude-mediated reciprocity process that elevates WTP at constant product quality, and that this reward is suppressed when effort is perceived as instrumentally motivated rather than genuinely invested. Because AI disclosure signals automated production, it frames the brand's communicative effort as cost-minimizing rather than creatively invested, thereby activating the suppression condition that eliminates the WTP premium that human effort would generate (Morales, 2005). In the high-end luxury market, the fundamental justification for a higher WTP stems directly from this perceived effort associated with the luxury craftsmanship (To et al., 2025; Xu & Mehta, 2022). When GenAI is employed, the resulting drop in perceived effort contradicts this justification. This cognitive devaluation deprives the luxury communication of its artisanal value, which subsequently reduces the financial premium the consumer is willing to allocate.

Methodologically, in the S-O-R framework, this internal cognitive evaluation process operates as a mediating variable because it is essential to explain why a true relationship exists between an exogenous construct (the Stimulus) and an endogenous construct (the Response), acting as the underlying mechanism that governs their association. By positioning Perceived Effort as a mediating variable, the model in Figure 3.1 aims to explain what happens in consumers' minds by demonstrating how the disclosure of an AI label in luxury advertising translates into diminished economic and psychological evaluations. The pathway from perceived effort suppression to reduced brand authenticity is further clarified by the distinction between indexical and iconic

authenticity established in Section 2.3.2. De Kerviler et al. (2022) demonstrate that indexical authenticity cues, grounded in verifiable production origin and genuine human creative investment, generate stronger brand ethicality and quality evaluations than iconic cues, which operate through aesthetic resemblance to established luxury conventions alone. Park & Ahn (2024) confirm that AI-generated luxury communications may achieve aesthetic parity with human equivalents, thereby satisfying iconic conditions, while simultaneously failing to meet the indexical requirements that luxury consumers rely upon to infer genuine creative origin. Consequently, the suppression of perceived effort induced by AI disclosure tends to undermine brand authenticity perceptions through the absence of indexical production cues, a process that may occur independently of the objective visual quality of the AI-generated stimulus.

Based on this theoretical rationale, the cognitive perception of low effort fundamentally undermines the luxury brands' value proposition. Therefore, Perceived Effort is proposed as the cognitive mechanism through which AI disclosure negatively influences consumers' economic and psychological evaluations of AI-generated luxury advertising. The following hypotheses formalize this mediating role:

H2: Perceived Effort mediates the negative effect of AI disclosure on consumers' Willingness to Pay and perceived Brand Authenticity, such that:

H2a: The disclosure of AI-generated visual content negatively impacts Perceived Effort compared to traditional human photography;

H2b: Perceived Effort positively impacts both consumers' Willingness to Pay and perceived Brand Authenticity.

3.2.3 The Emotional Pathway: Moral Disgust (Mediator 2)

The Organism (O) dimension of the S-O-R framework depicted in Figure 3.1 also captures a visceral affective response to the AI disclosure label, operating in parallel with and independently of the cognitive evaluation of effort. This affective response is theorized as moral disgust, a discrete negative emotion that arises in response to perceived violations of purity, authenticity, and moral integrity norms, which is functionally distinct from

anger and general negative affect (Giner-Sorolla & Chapman, 2017; Russell & Giner-Sorolla, 2011). Unlike anger, which is activated by harm and unfairness violations, moral disgust responds specifically to purity-related transgressions, including the perception that something genuine has been contaminated, replaced, or fraudulently represented (Giner-Sorolla & Chapman, 2017). In the context of luxury advertising, the disclosure of AI authorship may constitute precisely such a transgression: by revealing that a communication traditionally associated with human creative expression and artisanal investment was instead produced by a GenAI system, the disclosure signals a violation of the authenticity norms on which luxury brand identity is premised. In their study, Kirk & Givi (2025) demonstrate that AI-authored marketing communications elicit significantly elevated moral disgust relative to human-authored equivalents, establishing this affective response as a systematic consequence of AI disclosure.

The explicit disclosure of GenAI usage in luxury advertising has been shown to trigger emotional backlash. Recent research characterizes this phenomenon as the AI-authorship effect, demonstrating that when consumers recognize an emotional or creative communication was generated by an artificial entity rather than a human, it systematically evokes moral disgust (Kirk & Givi, 2025). This affective response arises because genuine communication requires congruence between a creator's internal state and its external expression (Kirk & Givi, 2025). Because AI systems lack genuine internal emotional states and human intentionality, their outputs in affective domains such as luxury advertising are evaluated by consumers as communicatively inauthentic, in that the communication cannot reflect a genuine emotional experience on the part of its author (Kirk & Givi, 2025). Empirical evidence from the study of Belanche et al. (2025) demonstrates that AI-generated images frequently evoke feelings of perceptual dissonance and are perceived as highly artificial and fake. Consequently, disclosing the use of GenAI may trigger an inference of manipulative intent on consumers, causing them to feel that the brand is trying to manipulate their perceptions through fabricated realities (Belanche et al., 2025).

Analogous to the psychological mechanisms whereby consumers experience moral disgust when they encounter a counterfeit or a fake luxury good, Kirk & Givi (2025) argue that explicitly labeling an advertisement as “AI-generated” violates the core expectations of human craftsmanship and authenticity, triggering a stronger moral disgust response. The mechanism through which moral disgust reduces brand authenticity perceptions operates through the psychological contagion dynamic documented by Amar et al. (2018), whereby moral impurity associated with a norm-violating communication spreads to contaminate the consumer's broader evaluation of the brand itself. Because the brand is perceived as the responsible communicative agent, the disgust elicited by an AI-generated luxury advertisement tends to extend beyond the specific stimulus, reducing authenticity perceptions even when the brand's heritage and product quality remain unchanged (Amar et al., 2018; Kirk & Givi, 2025). Romani et al. (2012) confirm that disgust-adjacent negative brand emotions produce measurable downstream penalties mediated by reductions in the evaluative integrity attributed to the brand, establishing moral disgust as the initiating event in a sequence through which the disclosure of AI authorship may degrade the consumer-brand relationship. Furthermore, as consumers actively seek to distance themselves from the source of the moral violation, this severe emotional reaction undermines the psychological justification required to sustain premium pricing, thereby significantly lowering their WTP (Kirk & Givi, 2025).

Methodologically, this affective reaction operates as a parallel mediating variable alongside Perceived Effort. As established by Hair et al. (2017), incorporating a mediator is essential to understand the underlying mechanism governing the association between the exogenous stimulus and the endogenous responses. By positioning Moral Disgust as an affective mediator, the model depicted in Figure 3.1 aims to demonstrate how the mandated disclosure of the AI label in a luxury advertising context translates into diminishing economic and psychological evaluations. Based on this theoretical rationale, the emotional reaction of moral disgust serves as a critical mechanism that undermines luxury brands' appeal and value. Therefore, Moral Disgust is proposed as the affective mechanism through which AI disclosure negatively influences consumers' economic and

psychological evaluations of luxury advertising. The following hypotheses formalize this mediating role:

H3: Moral Disgust mediates the negative effect of AI disclosure on consumers' Willingness to Pay and perceived Brand Authenticity, such that:

H3a: The disclosure of AI-generated visual content positively impacts Moral Disgust compared to traditional human photography;

H3b: Moral Disgust negatively impacts both consumers' Willingness to Pay and perceived Brand Authenticity.

3.2.4 The Moderating Role of the Need for Uniqueness

The cognitive and affective pathways described in Sections 3.2.2 and 3.2.3 establish the mechanisms through which AI disclosure negatively affects consumers' WTP and perceived Brand Authenticity. A fundamental limitation of treating these effects as uniform across all consumers, however, is that the magnitude of such responses may vary systematically as a function of individual-level traits, necessitating the identification of a boundary condition that accounts for this heterogeneity. The present study theorizes consumers' Need for Uniqueness as the relevant moderating variable, grounded in the premise that individuals differ systematically in the degree to which they are motivated to differentiate themselves from others through their consumption choices (Snyder & Fromkin, 1977; Tian et al., 2001).

Need for Uniqueness, as operationalized in the consumer behavior literature, refers to "an individual's pursuit of differentness relative to others that is achieved through the acquisition, utilization, and disposition of consumer goods for the purpose of developing and enhancing one's personal and social identity" (Tian et al., 2001). As established in Section 2.5.1, this disposition manifests behaviorally through the avoidance of similarity mechanism, whereby the perception that a previously distinctive brand association has become standardized or broadly accessible diminishes its symbolic differentiation value and suppresses consumer engagement (Tian et al., 2001). Luxury brands constitute a

particularly salient vehicle for the expression of this need, given that their scarcity, exclusivity, and dependence on human creative singularity provide the symbolic distinctiveness signals through which high-need-for-uniqueness consumers construct and communicate a differentiated identity relative to broader consumer groups (Belk, 1988; Simon & Fassnacht, 2019). Within this framework, the human creative origin of luxury advertising communications may function not merely as an aesthetic attribute but as an integral component of the exclusivity signal that sustains the identity investment of consumers whose luxury engagement is primarily motivated by differentiation (Loureiro et al., 2024; Tian et al., 2001).

The explicit disclosure that a luxury advertisement was generated by an AI embeds cues of automated, scalable production into the brand narrative, directly signaling the erosion of the human creative scarcity that sustains the distinctiveness function of luxury consumption. Loureiro et al. (2024) demonstrate that for consumers characterized by a high Need for Uniqueness, an AI-enabled experience must be genuinely distinctive to be positively evaluated, and that when the output is perceived as standardized, the avoidance of similarity mechanism is directly activated with downstream consequences for willingness to pay. Furthermore, as the widespread deployment of GenAI tools relies heavily on shared training data distributions, there is a systemic tendency toward output homogenization that reinforces the standardization signal embedded in the disclosure label (Wahid et al., 2025). For high-need-for-uniqueness consumers, this signal constitutes a direct threat to their social identity, depriving the communication of the human scarcity and craftsmanship typically required to differentiate their consumption choices from the mainstream (Loureiro et al., 2024).

Empirical support for this moderating logic is further provided by Granulo et al. (2021), who demonstrate that the preference for human over automated production is significantly stronger among consumers for whom distinctiveness constitutes a primary consumption goal, and that these consumers exhibit greater aversive responses to automation precisely because it eliminates the symbolic differentiation signal that justifies their premium investment. Since AI disclosure functions as precisely such a

standardization signal, its negative impact is expected to be systematically amplified among high-need-for-uniqueness consumers across both the cognitive and affective pathways established in Sections 3.2.2 and 3.2.3. The suppression of perceived creative effort and the inference of communicative inauthenticity are the specific cues through which the disclosure erodes the exclusivity signals that high-need-for-uniqueness consumers require from luxury brand communications, amplifying the negative consequences for both Perceived Effort and Moral Disgust relative to consumers for whom this distinctiveness requirement is less central (Giner-Sorolla & Chapman, 2017; Russell & Giner-Sorolla, 2011; Tian et al., 2001).

Methodologically, this study accounts for consumer heterogeneity by formalizing the Need for Uniqueness as a moderating variable in the S-O-R framework. As established by Hair et al. (2017), a moderation effect occurs when the strength or direction of a relationship between two constructs depends on the values of a third variable. By conceptualizing the Need for Uniqueness as a moderator in the S-O-R framework, this trait serves as a boundary condition that accounts for consumer heterogeneity, demonstrating that the effect of AI disclosure is not constant across all individuals, but instead it systematically influences their cognitive and emotional evaluations in a luxury context. Based on this theoretical rationale, Need for Uniqueness operates as a boundary condition that may intensify the negative psychological and economic evaluation of AI-generated luxury advertising. Since Need for Uniqueness is expected to intensify consumers' cognitive and emotional reactions to AI disclosure, it should also amplify the overall negative effect of the disclosure on WTP and Brand Authenticity. Therefore, the following hypotheses are formally proposed to test the moderating boundary condition of the S-O-R model:

H4: Consumers' Need for Uniqueness moderates the effect of AI disclosure on Perceived Effort and Moral Disgust, such that:

H4a: Consumers' Need for Uniqueness moderates the effect of AI disclosure on Perceived Effort, such that the negative effect of AI disclosure on Perceived Effort is stronger for consumers with higher levels of Need for Uniqueness.

H4b: Consumers' Need for Uniqueness moderates the effect of AI disclosure on Moral Disgust, such that the positive effect of AI disclosure on Moral Disgust is stronger for consumers with higher levels of Need for Uniqueness.

H4c: The negative indirect effect of AI disclosure on consumers' Willingness to Pay and perceived Brand Authenticity, mediated by Perceived Effort, is stronger for consumers with higher levels of Need for Uniqueness.

H4d: The negative indirect effect of AI disclosure on consumers' Willingness to Pay and perceived Brand Authenticity, mediated by Moral Disgust, is stronger for consumers with higher levels of Need for Uniqueness.

3.3 Summary of the Research Hypotheses

The preceding sections have formalized the theoretical relationships proposed in the dual-path S-O-R model, resulting in thirteen distinct hypotheses. At the core of this framework, Hypotheses H1a and H1b examine the total effects of the AI disclosure label on consumers' WTP and perceived Brand Authenticity, capturing the economic and psychological consequences of the stimulus. To better understand the internal processing within the Organism stage, Hypotheses H2 and H3 introduce two parallel mediating mechanisms, representing cognitive (Perceived Effort) and affective (Moral Disgust) responses to AI-generated content. Finally, Hypotheses H4a through H4d account for consumers' Need for Uniqueness as a boundary condition, moderating the initial impact of the AI disclosure label on the mediating mechanisms and, consequently, shaping the strength of the indirect effects on the outcome variables.

An overview of the complete hypothesis set is provided in Table 3.2. The following chapter builds on the established theoretical foundation by detailing the quantitative methodological design, the stimuli development process, and the specific analytical strategy employed to test the proposed relationships.

Table 3.2. Overview of the research hypotheses (Source: the author).

Hyp.	Statement
H1a	The disclosure of AI-generated visual content negatively impacts consumers' WTP compared to traditional human photography.
H1b	The disclosure of AI-generated visual content negatively impacts perceived Brand Authenticity compared to traditional human photography.
H2	Perceived Effort mediates the negative effect of AI disclosure on consumers' WTP and perceived Brand Authenticity.
H2a	The disclosure of AI-generated visual content negatively impacts Perceived Effort compared to traditional human photography.
H2b	Perceived Effort positively impacts both consumers' WTP and perceived Brand Authenticity.
H3	Moral Disgust mediates the negative effect of AI disclosure on consumers' WTP and perceived Brand Authenticity.
H3a	The disclosure of AI-generated visual content positively impacts Moral Disgust compared to traditional human photography.
H3b	Moral Disgust negatively impacts both consumers' WTP and perceived Brand Authenticity.
H4	Consumers' Need for Uniqueness moderates the effect of AI disclosure on the cognitive and emotional mediating pathways of the S-O-R model.
H4a	Consumers' Need for Uniqueness moderates the effect of AI disclosure on Perceived Effort, such that the negative effect of AI disclosure on Perceived Effort is stronger for consumers with higher levels of Need for Uniqueness.
H4b	Consumers' Need for Uniqueness moderates the effect of AI disclosure on Moral Disgust, such that the positive effect of AI disclosure on Moral Disgust is stronger for consumers with higher levels of Need for Uniqueness.
H4c	The negative indirect effect of AI disclosure on consumers' WTP and perceived Brand Authenticity, mediated by Perceived Effort, is stronger for consumers with higher levels of Need for Uniqueness.
H4d	The negative indirect effect of AI disclosure on consumers' WTP and perceived Brand Authenticity, mediated by Moral Disgust, is stronger for consumers with higher levels of Need for Uniqueness.

4. Methodology

Building upon the theoretical framework established in the preceding chapter, which proposed a Stimulus-Organism-Response (S-O-R) conceptual model, this chapter outlines the methodological paradigm employed to investigate how the disclosure of Generative AI (GenAI) in luxury advertising impacts consumers' cognitive evaluations, emotional reactions, and subsequent behavioral and psychological outcomes. The present chapter is structured as follows. First, Section 4.1 details the quantitative research approach and the explanatory, scenario-based, between-subjects experimental design. Section 4.2 delves into the stimuli development process, providing a comprehensive explanation of the state-of-the-art GenAI model utilized and the experimental manipulation of the disclosed label. Section 4.3 outlines the survey structure, the operationalization of the measurement scales, and the experimental validity of the study by detailing the implemented control and manipulation checks. Section 4.4 characterizes the sample population, alongside the execution of the data collection and cleaning processes adopted. Lastly, Section 4.5 introduces Partial Least Squares Structural Equation Modelling (PLS-SEM) as the chosen statistical framework for the data analysis.

4.1 Research Approach and Experimental Design

4.1.1 Quantitative Research Approach and Experimental Design

To empirically investigate the Stimulus-Organism-Response (S-O-R) framework proposed in Chapter 3, this study adopts a quantitative, deductive, and explanatory research paradigm. Unlike exploratory research, which is typically conducted when prior literature is limited, an explanatory study is mainly theory-driven and seeks to directly explain cause-and-effect relationships. This approach relies on a controlled experimental manipulation to identify the underlying psychological mechanisms, specifically the mediating and moderating variables, that drive consumer reactions toward AI-generated

advertisement disclosure. The research utilizes a cross-sectional, structured questionnaire to measure the impact of an AI label disclosure on consumers' Willingness to Pay (WTP) and perceived Brand Authenticity.

The empirical investigation is operationalized through a controlled, scenario-based, single-factor between-subjects online experiment with random assignment to conditions, where the individual consumer serves as the primary unit of analysis. To safeguard internal validity by controlling for pre-existing brand loyalty, equity, and familiarity, the experiment utilizes a fictitious luxury brand, "MERIDIAN," consistent with established methodological approaches (Ali et al., 2025; Kirk & Givi, 2025). The key study constructs are defined as follows: the Independent Variable (IV) is the authorship disclosure label in the luxury advertisement; the Dependent Variables (DVs) are consumers' WTP and perceived Brand Authenticity; capturing the underlying psychological mechanisms, the Mediating Variables (MVs) are defined as Perceived Effort (the cognitive pathway) and Moral Disgust (the emotional pathway); lastly, the Moderating Variable (Mod) is the consumers' Need for Uniqueness, which accounts for consumer heterogeneity.

Within this specific framework, respondents are randomly assigned to one of two experimental conditions: viewing either the MERIDIAN advertisement labelled as "Traditional Human photography" (the control group) or the exact same advertisement labelled as "Entirely generated by Artificial Intelligence" (the treatment group). This approach was a deliberate methodological choice to maximize internal validity. Specifically, as demonstrated by recent literature investigating AI-authorship and algorithmic aversion, employing an experiment where visual content is held constant isolates the psychological effects of the source disclosure (Heimstad et al., 2025; To et al., 2025). Using identical images controls for confounding visual variations and ensures that any observed change in the dependent variables is strictly related to the cognitive and affective reactions triggered by the AI label itself. Furthermore, restricting participants' exposure to a single condition prevents carryover effects and demand characteristics that are frequent in within-subjects designs, where respondents might consciously alter their evaluation upon recognizing the comparative nature of a study (Heimstad et al., 2025).

Finally, to analyze the complex relationships among these variables, the data analysis relies on Partial Least Squares Structural Equation Modelling (PLS-SEM). While PLS-SEM is traditionally recognized as highly suited for exploratory predictive modelling and theory development, its application in the present study is driven by its robustness in estimating complex, newly developed models and its capacity to maximize the explained variance in the dependent variables. Additionally, its robustness with smaller sample sizes and non-normal data distributions further justifies its application in testing the proposed experimental framework (Hair et al., 2017). Section 4.5 provides a detailed overview of PLS-SEM application within this study.

4.1.2 Declaration of AI Tools Utilization

In compliance with institutional guidelines regarding the responsible use of GenAI in academic work issued by the Ca' Foscari University of Venice, the present section provides a transparent and detailed account of the AI tools employed and the purpose for which they were applied. All outputs produced with the assistance of AI tools were subject to verification by the author, to ensure their accuracy, quality, and alignment with the present study.

Three AI-based tools were employed across distinct phases of the research process. First, regarding the generation of the visual stimuli, FLUX.2 ([dev] version) was utilized to produce the AI-generated luxury advertisement image used in the experimental manipulation. Given the importance of this tool to the study design, a comprehensive discussion of its architecture, selection rationale, and deployment is reserved for Section 4.2. Second, Gemini Pro, accessed through the institutional university account in accordance with the agreements stipulated between the University and the service provider, was employed exclusively during the preliminary phase of the literature search documented in Section 2.1. Its application was limited to the construction and refinement of keyword combinations used to query the selected academic databases. At no stage was Gemini Pro or any other AI-tool employed to read, summarize, or synthesize journal articles and other sources. The full-text review and critical evaluation of all primary and secondary literature were conducted exclusively by the author. Third, Claude (Anthropic,

Pro subscription) was employed in two ways throughout the thesis. The first concerned writing quality: Claude was used to identify grammatical errors, improve paragraph clarity, and highlight points requiring further development. This function was applied across all chapters. The second concerned analytical support: Claude was used to assist in the interpretation of the PLS-SEM results presented in Chapter 5, supporting the contextualization of statistical outputs within the study's theoretical framework. All AI-assisted content was critically reviewed and verified by the author before inclusion in the final text.

4.2 The Stimuli Development and the Image Generation Process

4.2.1 The Generative AI Model (FLUX.2 [dev])

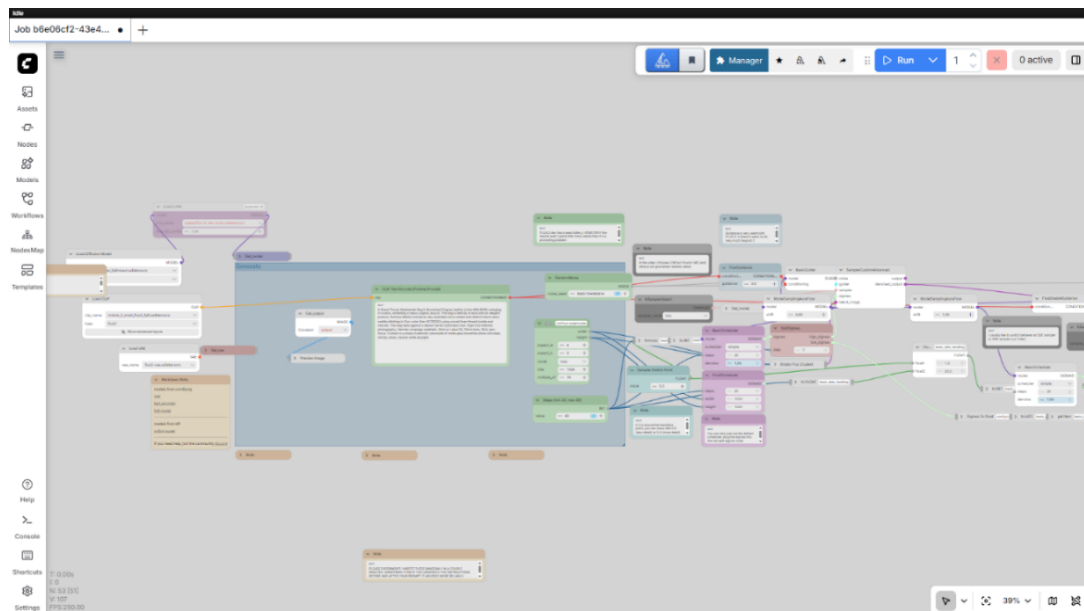
The recent advancements of GenAI, particularly of Large Language Models (LLMs) and text-to-image diffusion models (such as OpenAI's DALL-E 3 or Midjourney), have transformed marketing content creation (Prasanna & Kushwaha, 2025b). However, the integration of AI-generated content in high-involvement contexts, such as the luxury sector, requires careful consideration. Conventional GenAI models usually struggle with complex visual details, such as accurately managing the lighting physics, and generating human features, particularly hands (Borji, 2023; Hartmann et al., 2025). These visual anomalies are usually referred to as "AI artifacts" and can trigger the so-called uncanny valley effect, as established in Chapter 2 (Kirk & Givi, 2025). This effect is a phenomenon in which an artificial entity appears somewhat unnatural, eliciting feelings of unease and immediately signaling to consumers that the image is fake (Hartmann et al., 2025; Mori, 1970). For luxury brands, these visual "AI artifacts" destroy the perception of premium value, and diminish perceived brand authenticity (To et al., 2025). Therefore, the successful integration of AI-generated images in hedonic and high-involvement contexts, such as the luxury sector, has been noticeably slow because this market requires AI models capable of producing flawless visuals that adhere strictly to specific brand values and guidelines (Hermann & Puntoni, 2025). To achieve the level of quality required to test

the previously proposed framework, this study utilizes FLUX.2 (specifically the [dev] version), an open-weight, state-of-the-art image generation model released by Black Forest Labs in November 2025 (Black Forest Labs, 2025a). Compared to more traditional and consumer models, FLUX.2 relies on an advanced architecture that provides superior semantic understanding of instructions and accurately accounts for details in the image generation process, thus minimizing the occurrence of AI artifacts while producing photorealistic images (Black Forest Labs, 2025b).

In the context of this study, employing a model capable of generating the highest level of realism was key. To accurately test the impact of the AI authorship disclosure label (the independent variable), the visual stimulus had to exhibit the premium-quality characteristics typical of luxury brands. If the generated advertisement exhibited noticeable AI artifacts or low visual appeal, participants' responses could be driven by the execution quality of the image or by their recognition of it as AI-generated, rather than by the AI authorship label disclosure itself (Heimstad et al., 2025). Therefore, by generating a highly realistic output, the application of FLUX.2 ensures that any observed changes in consumers' WTP or perceived Brand Authenticity were strictly attributable to the cognitive and affective reactions triggered by the AI label (Kirk & Givi, 2025; To et al., 2025). Furthermore, the selection of the FLUX.2 [dev] variant was driven by two factors. First, it was the only model available at the time of the experiment, released alongside the smaller [klein] version. Second, its open-weight nature and non-commercial, research-focused license provided a significant advantage over competing GenAI models (Black Forest Labs, 2025a). Traditional state-of-the-art models, such as OpenAI's DALL-E 3 or Midjourney, operate as closed-source "black boxes," accessible exclusively through paid subscriptions. These closed ecosystems impose continuous financial fees and strict usage limits. Conversely, FLUX.2 [dev] is entirely free for research purposes, eliminating these recurring costs, and providing open access since it can be deployed completely offline on local hardware architectures (Black Forest Labs, 2025b). Additionally, the study's implementation of explicit AI labeling directly aligns with emerging regulatory frameworks, such as the European Union AI Act (Article 50), which formalizes transparency obligations requiring creators and brands to clearly disclose when content is artificially generated (European Parliament, 2024b; To et al., 2025).

The image generation process of the present study was executed offline. As briefly discussed, traditional commercial models operate as closed-source systems, and a key limitation is that the same text prompt may yield entirely different images across generations. This lack of reproducibility occurs because these systems do not allow users to save the exact starting conditions of a generated image. Therefore, this process inherently involves randomness and the exact same text prompt will produce a different image at each iteration (Cillo & Rubera, 2025). Conversely, FLUX.2 [dev] allowed for a fully offline and free deployment. The model was run within ComfyUI, an advanced, open-source, node-based graphical interface designed specifically for complex GenAI workflows (Comfy, 2025a, 2025b). ComfyUI allows researchers to construct visual pipelines by linking modular blocks across a canvas, and the process was executed on a dedicated local computational environment equipped with an NVIDIA RTX 4090 GPU (24GB of VRAM). This setup granted full control over the image generation process, ensuring that the image creation remained transparent and reproducible. An example of the workflow created for this study is depicted in Figure 4.1 using the standard available workflow that opens when FLUX.2 [dev] is loaded.

Figure 4.1. ComfyUI workflow used to create the GenAI image (Source: the author).



Building upon these technical foundations, Section 4.2.2 will detail the specific prompt engineering, the parameter configurations, and the steps executed within ComfyUI to generate the final visual stimuli adopted in the questionnaire.

4.2.2 The Image Creation Process

The main product selected for the brand MERIDIAN was a “Weekender Bag.” This product was chosen to ensure the internal validity of the experiment and prevent confounding demographic biases. In experiments, utilizing a commonly gendered product, such as a woman’s purse or a man’s watch, could trigger unwanted gender-based biases, varying levels of product involvement, or different levels of product knowledge among respondents. Therefore, a weekender bag serves as a unisex accessory. This selection ensured a gender-neutral stimulus and standardized the evaluation context for all participants (Heitmann et al., 2025).

To generate the visual stimulus, the text prompt was engineered in strict adherence to the official FLUX.2 documentation provided by Black Forest Labs (Black Forest Labs, 2025c). Since FLUX.2 does not utilize negative prompts (which usually state what an image should not contain), the instructions required highly specific visual constraints structured around the subject, action, style, and context. Moreover, because FLUX.2 features an advanced semantic understanding of spatial depth, lighting, and textures, it requires highly detailed input to maximize its photorealistic capabilities (Black Forest Labs, 2025b). Google’s LLM, Gemini Pro, was utilized in the input creation process to refine the technical photographic terminology, without influencing the content creation or design thinking process, as done in the study of Prasanna & Kushwaha (2025b). This mixed (human+AI) approach ensured that the prompt contained accurate specifications for the camera, lenses, and professional studio lighting setups. The full prompt used in the present study is as follows:

A Grand Tourer Weekender Bag in Burnished Cognac leather (color #8A4B38) swinging in motion, exhibiting a heavy organic slouch. The bag is held by a hand with an elegant posture. Surface details include an oily, hydrated micro-sheen and distinct hand-sewn saddle

stitching in Ecru color (hex #C2B280) using waxed linen thread (matte and natural). The bag rests against a dense Camel Cashmere Coat. High-End Editorial photography, Hermès campaign aesthetic. Shot on Leica S3, 70mm lens, f/8.0, pan-focus. Context is a sharp modernist colonnade of matte grey travertine stone and steel, strictly urban, neutral white daylight.

Within the ComfyUI environment, the generation parameters were configured to balance high-quality output with the computational feasibility of the researcher's local hardware (NVIDIA RTX 4090). The exact model checkpoint utilized was flux2_dev_fp8mixed.safetensors. While the model is capable of producing ultra-high-resolution images, the output resolution for this study was standardized at 1024x1024 pixels. This resolution was deemed optimal for the experimental survey format and required approximately four minutes of generation time per image, avoiding the operationally unfeasible eight-minute processing time required for standard 1920x1080 resolutions. Furthermore, this resolution aligns directly with the standard image format of Instagram's square post format (1:1 aspect ratio), which represents one of the main digital channels for advertising (Hartmann et al., 2025). Presenting the stimulus at a resolution consistent with social media advertising ensures that participants are exposed to a format they would realistically encounter in an authentic brand communication context. All generation parameters were fixed across iterations. The seed was varied during the exploratory phase to generate multiple candidate images, while the final selected image was generated using a fixed seed (65851154065814) to ensure full reproducibility. The sampler type was set to Euler, and the generation sequence was set to 40 steps. FLUX.2 uses a specific guidance parameter (CFG) to determine how strictly the AI should follow the text prompt versus allowing for creative deviation. In the present study, the guidance scale was locked at 3.0 to ensure strong adherence to the prompt (Black Forest Labs, 2025b). To increase the level of details, an upscaling technique was used within FLUX.2. Specifically, the first 28 steps were dedicated to generating the core image, while the remaining 12 steps functioned as a detailing upscaler, enhancing the level of detail in the image without changing its structure. All remaining model

sufficient image diversity and maintaining the operational constraints of a real-world work environment. The selection process was guided by strict exclusion and inclusion criteria. First, any images exhibiting noticeable "AI artifacts," such as structural deformities in the bag's handles or human hands, illogical lighting shadows, or floating elements, were immediately discarded (Hartmann et al., 2025). Figure 4.4 depicts an example of rejected images. The final image was selected on the basis of its high visual realism, absence of AI artifacts, and its faithful adherence to the luxury brand aesthetic specified in the prompt. To validate the selected stimulus, a pre-test was conducted among a small group of peer graduate students ($N = 5$), confirming that the selected image successfully evoked the intended perceptions of high realism and luxury without disclosing the intentions of the study. The finalized, unaltered base image is depicted in Figure 4.5. For the experimental deployment in the survey, the base image was subjected to minor post-production editing using the graphic design platform Canva to apply the authorship disclosure labels required for the experimental conditions. The composition, lighting, camera angle, and background were held constant across all conditions. Only the label (AI vs. human) was manipulated, meaning no visual differences existed between conditions.

Figure 4.4. Examples of rejected images with AI Artifacts (Source: FLUX.2 [dev]).



Figure 4.6a depicts the stimulus (in Italian) with the "MERIDIAN Campaign S/S 2025. Photography by G. Von Hassel" label (the control group), while Figure 4.6b depicts the exact same image featuring the "MERIDIAN Campaign S/S 2025. Image entirely generated with Artificial Intelligence" label (the treatment group). The attribution of the control stimulus to a fictitious photographer was a deliberate choice to maintain the realism of a luxury campaign, while preventing any pre-existing consumer biases associated with real-world photographers.

Figure 4.5. *The final, unaltered image, selected for this study (Source: FLUX.2 [dev]).*



This rigorous control over the image generation process ensures that any observed effects can be attributed solely to the AI disclosure manipulation, thereby directly supporting the internal validity of the hypothesized relationships.

Figure 4.6a (left). The stimulus depicting the Human Label (Source: FLUX.2 [dev]).

Figure 4.6b (right). The stimulus depicting the GenAI Label (Source: FLUX.2 [dev]).



4.3 Questionnaire Structure and Measurement Scales

4.3.1 Questionnaire Structure

To collect the primary data, a cross-sectional, structured online questionnaire was developed. The survey was administered using Qualtrics, which is widely recognized in academic research for its advanced flow capabilities. Compared to other survey platforms such as Google Forms, Qualtrics allows for the automated and randomized assignment of participants to distinct experimental conditions while facilitating data quality controls. Because the target sample primarily consisted of Italian participants, the questionnaire was administered in Italian. To ensure linguistic equivalence and preserve the precise semantic meaning of the original English measurement scales, a rigorous forward-and-back translation procedure was implemented prior to the questionnaire publication. The survey followed a logical sequence designed to guide respondents step-by-step, beginning and ending with simpler questions. The estimated completion time of the questionnaire was between four and five minutes.

The survey began with an introductory section outlining the general purpose of the research, assuring data anonymity, and securing informed consent. The introduction described the study in broad terms, referring to the analysis of consumer perceptions of a luxury advertisement without disclosing the specific focus on the authorship of the image (AI vs Human). Immediately after, participants were asked to answer a preliminary control question to establish their baseline interest in luxury. Measuring this baseline interest prior to any experimental exposure ensured that pre-existing category involvement could be controlled. Following this introductory stage, the experimental manipulation phase began. Utilizing the randomization logic embedded within Qualtrics, respondents were randomly assigned to one of two conditions and exposed to the same visual advertisement but with a different label. Here, the advertisement was accompanied by either a human-authorship disclosure label (the control group) or an AI-generated disclosure label (the treatment group) as shown in Figure 4.6. Furthermore, participants were required to view the image for a minimum of 20 seconds. A page-level timer prevented progression until this exposure threshold had elapsed.

Following the stimulus exposure, the questionnaire presented to participants the core evaluation section. Respondents were asked to evaluate the Visual Appeal of the advertisement they just saw to empirically verify the attractiveness of the stimulus itself across both conditions (to assess whether the label disclosure influenced participants' evaluations of the advertisement). The survey then proceeded to verify the mediating variables, by sequentially capturing the cognitive and emotional psychological pathways triggered by the advertisement, specifically the Perceived Effort and the Moral Disgust. Next, participants evaluated the primary dependent variables, assessing their perception of Brand Authenticity and their monetary Willingness to Pay (WTP) for the product depicted in the advertisement. This section was then followed by the measurement of the moderating variable, the Need for Uniqueness, to account for individual consumer heterogeneity. Lastly, the questionnaire was concluded by collecting standard demographic information about gender, age, nationality, education level, and employment status. Positioning these simpler, descriptive demographic questions at the very end of the survey helped minimize cognitive fatigue during the more demanding core experimental evaluations.

Prior to launching the questionnaire publicly, a small pre-test ($N = 5$) was conducted to identify any initial structural flaws and to ensure a smooth and logical survey flow. The primary purpose of this pilot evaluation was to assess the comprehensibility of the questions, refine the wording, and verify the accuracy and consistency of the Italian translation. Based on the qualitative feedback provided by these initial participants, minor revisions were implemented to optimize language comprehension and finalize the survey. Furthermore, the pre-test also revealed a key issue: the AI or Human label was not sufficiently visible within the stimulus. This limitation affected the overall clarity of the experimental manipulation. To address this, the label was enlarged within the image and also explicitly included in the question text during the stimulus exposure stage. The final questionnaire was distributed via an anonymous online link generated by the Qualtrics platform. The data collection phase started on February 28, 2026, and officially concluded on March 10, 2026, after three days with no more answers. The survey was actively shared through social media and messaging applications, specifically Instagram and WhatsApp. Participants were initially recruited through the researcher's personal network and in strict compliance with GDPR and academic ethical standards. All participants provided informed consent, retained the right to withdraw at any time, and participated entirely anonymously without the collection of any personally identifiable data. To expand the reach of the questionnaire, these initial respondents were explicitly encouraged to forward the survey link to two or three of their peers.

Methodologically, this strategy constitutes a non-probability sampling approach known as snowball sampling (Atkinson & Flint, 2001). In this process, the researcher identifies an initial group of participants who then recruit new subjects from their own social networks, allowing the sample size to grow organically. This technique was deliberately selected for its practical advantages in online survey-based research. It relies on established interpersonal connections, which facilitate a sense of trust among potential respondents, making them more comfortable participating in the survey. Furthermore, it offers a highly efficient and cost-effective means to rapidly diffuse the survey link and gather diverse data within a relatively short timeframe (Atkinson & Flint, 2001). It must be acknowledged that snowball sampling inherently carries the risk of selection bias, as individuals tend to recruit peers with similar demographic or psychographic traits within

their own social circles. This may reduce the overall representativeness of the sample and potentially limit the generalizability of the findings. Nevertheless, this approach remains highly suitable for explanatory quantitative studies. The following section (Section 4.3.2) details the specific multi-item measurement scales adapted from the academic literature. Subsequently, to guarantee the internal validity of the experimental design and the overall quality of the respondent data, Section 4.3.3 outlines the implementation of the manipulation and attention checks.

4.3.2 Measurement Scales

In the present study, the variables under investigation are abstract, complex, and unobservable, and therefore function as latent constructs. Since such constructs cannot be directly measured, they must be assessed indirectly through observable indicators that act as proxy variables (Hair et al., 2017). Each measurement item captures a specific dimension of a broader concept, thereby enabling its empirical operationalization. To accurately capture the multidimensional nature of these constructs, this study employs multivariate measurement through the use of multi-item scales, with the exception of the control variable Luxury Interest and the single open-ended WTP. The adoption of multiple indicators enhances overall measurement accuracy by reducing the random measurement error, defined as the difference between a construct's true underlying value and the value obtained during the measurement process (Hair et al., 2017).

To ensure high reliability and construct validity, all measurement scales were derived from previously validated academic literature. Minor wording adaptations were introduced to align the items with the specific experimental context, namely, the evaluation of an AI-generated versus a human-photographed labelled luxury advertisement. Furthermore, to ensure consistency and reduce respondents' cognitive fatigue, all latent constructs were measured using a seven-point scale format. Although certain scales originally employed five-point formats, they were adapted to a seven-point scale to ensure consistency across all constructs. The full set of measurement scales adopted in this study is reported in Table 4.1, resulting in 17 items used to measure 7 distinct constructs, in addition to the experimentally manipulated independent variable.

Table 4.1. *Measurement Scales (Source: the author).*

Construct	Item	Source
AI Disclosure	Condition A: "Photography by G. Von Hassel"	Jin & Tao (2025)
	Condition B: "Image entirely generated by Artificial Intelligence"	Binary (0 = "Human"; 1 = "AI")
Luxury Interest	How interested are you in the world of luxury products and accessories?	Zaichkowsky (1985); Bearden et al. (2011) Seven-point Likert scale (1 = "Not at all," and 7 = "Extremely")
Visual Appeal	1. The way the product is displayed in the image is attractive.	Mathwick et al. (2001); Bearden et al. (2011) Seven-point Likert scale (1 = "Strongly disagree" and 7 = "Strongly agree")
	2. The image is aesthetically appealing.	
	3. I like the way the image looks.	
Perceived Effort	1. Time required to create the image (1 = Very little time, 7 = A great deal of time)	De Kerviler et al. (2022); To et al. (2025) Seven-point Semantic Differential scale (item-specific anchors)
	2. Effort required to create the image (1 = Very little effort, 7 = A great deal of effort)	
	3. Difficulty of creating the image (1 = Very easy, 7 = Very difficult)	
Moral Disgust	1. To what extent does the way this image was created make you feel disgusted?	Russell & Giner-Sorolla (2011); Kirk & Givi (2025) Seven-point Likert scale (1 = "Not at all," and 7 = "Very much")
	2. To what extent does the way this image was created make you feel repulsed?	
	3. To what extent does the way this image was created make you feel sickened?	

Brand Authenticity	1. How authentic do you consider the brand to be?	Luffarelli et al. (2019); Bruner (2021) Seven-point Likert scale (1 = "Not at all," and 7 = "Extremely")
	2. How trustworthy do you consider the brand to be?	
	3. How credible do you consider the brand to be?	
Willingness to Pay	What is the MAXIMUM price you would be willing to pay to purchase this bag?	Jones (1975); Fuchs et al. (2010) Open-ended question (maximum willingness to pay in euros)
Need for Uniqueness (Avoidance of Similarity dimension)	1. I avoid products bought by the average consumer.	Tian et al. (2001); Bearden et al. (2011) Seven-point Likert scale (1 = "Strongly disagree," and 7 = "Strongly agree")
	2. I lose interest in brands that become too popular.	
	3. The more common a product, the less I want it.	

The independent variable, AI disclosure, was manipulated as a binary categorical value (0/1), by exposing participants to either a human-authorship condition ("Photography by G. Von Hassel") or an AI-authorship condition ("Image entirely generated by Artificial Intelligence"), an approach adapted from Jin & Tao (2025).

To isolate the effects of this experimental manipulation, two control variables were established. First, to account for pre-existing category involvement, Luxury Interest was measured utilizing a single-item, seven-point Likert scale ranging from 1 = "Not at all" to 7 = "Extremely," evaluating respondents' baseline interest in the world of luxury products and accessories. A single-item measure was adopted based on the assumption that interest in luxury products represents a concrete and straightforward concept that can be reliably assessed with a single, direct indicator. This construct has been adapted from Zaichkowsky (1985) and Bearden et al. (2011). Second, Visual Appeal was measured to verify that the stimulus was perceived as attractive regardless of the authorship label.

This construct captures the aesthetic appreciation, design quality, and physical attractiveness inherent in a visual advertisement. It was assessed using a three-item, seven-point Likert scale ranging from 1 = "Strongly disagree" to 7 = "Strongly agree," by asking participants to evaluate the extent to which the product display was attractive, aesthetically appealing, and pleasing to look at, adapting the scale from Mathwick et al. (2001) and Bearden et al. (2011). Measuring Visual Appeal across both conditions serves a dual purpose: it allows the perceived attractiveness of the stimulus to be monitored across the authorship label conditions, and it allows the variable to be controlled for within the structural model, ensuring that any observed effects on the dependent variables are attributable to the AI disclosure label rather than to differences in aesthetic appeal.

To capture the underlying psychological mechanisms driving consumer responses, two mediating variables were established, namely Perceived Effort and Moral Disgust. The former represents the cognitive pathway and acts as a psychological signal of commitment, dedication, and the level of physical labour and skill employed during the production process. This construct was adapted from De Kerviler et al. (2022) and To et al. (2025) and helps to capture the effort heuristic. This was measured using a three-item, seven-point semantic differential scale (with item-specific anchors reflecting low versus high levels of time, effort, and difficulty) asking participants to describe the advertisement's creation process by rating it as time-intensive, effortful, and difficult to create. The latter variable, Moral Disgust, operates in parallel to Perceived Effort as the emotional pathway and was adapted from Kirk & Givi (2025) and Russell & Giner-Sorolla (2011), and evaluates the intensity of consumer's emotional repulsion. It was measured using a three-item, seven-point Likert scale ranging from 1 = "Not at all" to 7 = "Very much," assessing the extent to which the image creation process made respondents feel disgusted, repulsed, and sickened.

The Response (R) part of the Stimulus-Organism-Response (S-O-R) framework was captured by two dependent variables. Here, the perceived Brand Authenticity evaluates the subjective genuineness ascribed to a brand, determining whether the brand stays true to its promises, is intrinsically motivated, and avoids misrepresentation. The scale was

adapted from Luffarelli et al. (2019) and Bruner (2021), and was measured using a three-item, seven-point Likert scale ranging from 1 = "Not at all" to 7 = "Extremely," evaluating the extent to which the brand appeared authentic, trustworthy, and credible. Alongside this psychological outcome, the economic response was captured through the WTP. To prevent potential reference-price anchoring effects, WTP was measured using a direct, open-ended question requiring participants to state the maximum price in euros they would pay to purchase the advertised luxury weekender bag, an approach consistent with established literature and adapted from Jones (1975) and Fuchs et al. (2010).

Lastly, to account for individual consumer heterogeneity, the moderating variable, consumers' Need for Uniqueness, was operationalized. This trait reflects an individual's motivation to pursue difference relative to others through the acquisition and use of consumer goods, to enhance both their self-image and their social image (Tian et al., 2001). The scale was adapted from Tian et al. (2001) and Bearden et al. (2011), and it specifically captures the Avoidance of Similarity dimension of need for uniqueness. The construct was assessed using a three-item, seven-point Likert scale ranging from 1 = "Strongly disagree" to 7 = "Strongly agree," capturing the respondents' tendency to lose interest in popular brands, avoid commonplace products, and reject goods widely accepted by the average consumer.

All latent variables (apart from WTP) were measured utilizing seven-point scales (Likert and semantic differential). This was a deliberate decision made to offer respondents a wider spectrum of choice, thereby facilitating a more precise capture of their subjective evaluations. Furthermore, to ensure consistency and straightforward interpretation during the analysis, all items across these scales were positively coded, so that higher numerical values represented stronger agreement or a higher level of the underlying construct. In Structural Equation Modelling (SEM), it is crucial to carefully code ordinal scales to approximate the assumption of equidistance between response categories (Hair et al., 2017). To achieve this, the scales were designed to present strict structural symmetry around a middle category utilizing clearly defined linguistic qualifiers for each point. When a Likert scale is perceived as both symmetric and equidistant, it allows the

ordinal scale to effectively behave as interval-level measurement for the data analysis. Furthermore, because the survey was administered to respondents in Italian, maintaining methodological transparency is paramount. To guarantee complete comparability between the two linguistic versions, the full set of measurement items summarized in Table 4.1 has been translated into Italian and is provided for reference in Appendix A. Finally, adapting these measurement scales from established literature published in top-tier academic journals (ABS 3 or higher ranked journals) ensured content validity for the study. The statistical quality and specification of the measurement models is empirically assessed in Chapter 5.

4.3.3 Attention and Manipulation Checks

In quantitative, self-administered online experiments, it is important to implement a rigorous control mechanism to ensure data quality and confirm that participants are reading the instructions carefully rather than responding arbitrarily. To identify low-effort respondents, a single attention check was put in place within the measurement scales, and was positioned in the central part of the questionnaire. This attention check functioned as a filter for inattentive participants, by explicitly stating to them that they had to select a predetermined response option. Any participant who failed to comply with this instruction was then removed from the sample, as their inclusion would introduce substantial measurement error and compromise the overall reliability of the dataset. Beyond a general attention mechanism, it was important to objectively verify that the experimental manipulation was successfully understood by respondents. To achieve this, two memory-based manipulation checks were placed after the dependent variables questions (Brand Authenticity and WTP). This was a deliberate choice to avoid prematurely revealing the true research objective to participants while they were evaluating the dependent variables.

The first manipulation check was created to verify the success of the independent variable manipulation regarding the AI-disclosure label. If a participant could not correctly recall the authorship label they were exposed to, the experimental manipulation failed to register in participants' awareness, indicating a lack of manipulation awareness.

Consequently, their subsequent evaluations could not be validly used to measure the psychological or behavioral effects of the disclosure. The second manipulation check was implemented to verify that participants had actively processed the visual content of the advertisement, and had read the various questions consciously by asking them to recall the correct product category shown in the image. The exact constructs, items, and corresponding response options utilized for these control and manipulation measures are detailed in Table 4.2. Because the survey was administered in Italian, the exact translated wording used for these checks can be found in Appendix B.

Table 4.2. *Summary of Attention and Manipulation Checks (Source: the author).*

Control Measure	Item	Response Option
Attention Check	Please select the option “4 – Neither agree nor disagree” from the list below.	Seven-point Likert scale (1 = “Strongly disagree,” and 7 = “Strongly agree”) (Correct response required: Option 4)
Manipulation Check 1	Do you remember how the advertising campaign you evaluated at the beginning of the questionnaire was created?	1. Traditional photography 2. Image entirely generated by Artificial Intelligence (Correct response depends on assigned condition)
Manipulation Check 2	Besides the production technique, do you remember which product category was shown in the advertising image?	1. A weekender bag 2. A wristwatch (Correct response required: Option 1)

Ultimately, to preserve the internal validity and the causal integrity of the experiment, an exclusion procedure was put in place. Any individual who failed to respond correctly to even one of these control measures was deemed unreliable for the purposes of hypothesis testing. These respondents were systematically screened out and excluded from the final analysis, guaranteeing a final valid sample of highly attentive participants, upon whom the experimental manipulation had successfully operated.

4.4 Sample and Data Collection

4.4.1 Target Population and Sampling Strategy

The data collection phase was conducted using a structured, self-administered online questionnaire, developed and distributed via an anonymous link generated by the Qualtrics platform. The target population of the present study comprised adult Italian consumers, as they are more likely to possess the purchasing power required for luxury products, with the individual consumer serving as the unit of analysis. This decision was also based on the expectation that the majority of respondents would be Italian speakers. Following a small pre-test (N=5), described in Section 4.3.1, the data collection period officially started on February 28, 2026, and concluded on March 10, 2026, following three consecutive days of inactivity. To reach the intended audience, the survey link was actively distributed through social media platforms and messaging applications, specifically Instagram and WhatsApp. Initial recruitment was conducted through the researcher's personal network, after which participants were explicitly encouraged to forward the survey link to their peers and acquaintances to expand the study's reach.

Upon accessing the survey, participants were randomly assigned to one of two experimental conditions, exposing them to either an AI-generated or a human-created luxury advertisement. This randomized between-subjects procedure was designed to distribute participants evenly across the groups and was implemented through the Qualtrics randomizer function. The distribution technique utilized to reach these participants is snowball sampling, a technique where the researcher identifies an initial

group of participants who then recruit new subjects from their own social networks, allowing the sample size to grow organically. Snowball sampling advantages and problems were detailed in Section 4.3.1 (Atkinson & Flint, 2001). Ultimately, at the conclusion of the data collection period, an initial pool of 202 raw responses was recorded.

4.4.2 Data Preparation and Exclusion Criteria

The data screening process began with an initial pool of 202 raw responses collected at the conclusion of the survey period. To ensure the integrity and reliability of the dataset, a preliminary data screening was conducted to identify and remove incomplete submissions. Specifically, 25 responses were removed because participants abandoned the questionnaire before completion. Within this group, nine individuals did not progress past the initial access screen, ten dropped out during the stimulus exposure phase, four completed only part of the survey during the measurement section, one respondent exited after the WTP question, and one exited at the final demographic section. In accordance with methodological best practices suggested by Hair et al. (2017) for handling missing data, a strict listwise deletion method was applied. This ensures that only cases with complete data across all measured variables are retained, yielding a preliminary pool of 177 fully completed surveys. Subsequently, the dataset was screened for abnormally rapid completion times. An analysis of the response durations confirmed that no speeders were present, as all participants completed the questionnaire within a timeframe ranging from 3 to 9 minutes.

Following the removal of incomplete responses, the dataset was evaluated using the attention check to identify careless or low-effort responders. Participants were explicitly instructed to select the "4 - Neither agree nor disagree" option on a seven-point Likert scale. Demonstrating a high level of cognitive engagement with the survey, 100% of the 177 remaining respondents successfully passed this attention check, confirming their attentiveness during the evaluation of the core measurement scales. Furthermore, two memory-based manipulation checks were administered within the survey. These checks were necessary to verify that the experimental manipulation had successfully registered

in participants' conscious awareness and to preserve the internal validity of the experiment. First, participants were tested on their recall of the authorship label assigned to their specific condition (Human versus AI). A total of 21 respondents failed to accurately recall the label they were exposed to and were consequently excluded. Second, stimulus integrity was assessed by asking participants to identify the product category displayed in the advertisement. An additional 3 participants failed to correctly identify the weekender bag and were removed, leaving a total of 153 valid responses.

In the final stage of data preparation, the remaining observations were screened for extreme outliers and protest responses, specifically focusing on the open-ended WTP variable. As highlighted by Hair et al. (2017), retaining implausible extreme values in PLS-SEM can severely skew the mean, artificially inflate variances, and ultimately distort the parameter estimates. Consequently, seven respondents were removed at this stage: one respondent was removed for entering an extreme and implausible WTP value of €16,000 for the weekender bag, and six respondents were excluded for providing protest responses. These protest responses, characterized by clearly unrealistic values such as €0, €1, or €20 for a high-end luxury accessory, were identified as low-effort responses and excluded. Conversely, low but theoretically plausible values (e.g., €50 to €100) were retained, as these economic evaluations could legitimately reflect the respondent's true assessment of the product, particularly within the AI-generated stimulus condition. Additionally, a final screening conducted by computing the intra-respondent standard deviation (SD) across the fifteen reflective scale items, identified one additional respondent exhibiting a low SD ($SD = 0.46$). This case was consequently excluded. Upon the completion of the screening and cleaning procedures, a final reliable sample of 145 valid responses was obtained for the statistical analysis. Table 4.3 provides a comprehensive summary of the data collection and screening process.

Following the data cleaning process, the final sample exhibited a slight imbalance across the two experimental conditions: the human-photography control condition comprised 65 participants, while the AI-generated treatment group was made up of 80 participants. This variance resulted from the data exclusion and cleaning process, particularly

regarding the manipulation checks. However, this unequal allocation does not compromise the validity of the analysis. According to Hair et al. (2017), PLS-SEM is robust to unequal group sizes, and both subsamples comfortably exceed the minimum threshold required for reliable parameter estimation, as detailed in Section 4.5.

Table 4.3. Summary of Data Cleaning and Exclusion Process (Source: the author).

Survey Section	Stage Description	N. of Respondents
Survey Introduction	Total questionnaire opened	202
	<i>Drop-out in this section</i>	(9)
Stimulus Stage	Total n. of respondents	193
	<i>Drop-out in this section</i>	(10)
Variable Measurements Stage	Total n. of respondents	183
	<i>Drop-out in this section</i>	(4)
WTP Assessment Stage	Total n. of respondents	179
	<i>Drop-out in this section</i>	(1)
Demographics Stage	Total n. of respondents	178
	<i>Drop-out in this section</i>	(1)
Rapid Completion Time	Total n. of respondents	177
	<i>Drop-out in this section</i>	(0)
Total n. of completed responses		177
Attention Check	Total n. of respondents	177
	<i>Drop-out in this section</i>	(0)
Manipulation Check 1	Total n. of respondents	177
	<i>Drop-out in this section</i>	(21)

Manipulation Check 2	Total n. of respondents	156
	<i>Drop-out in this section</i>	(3)
WTP Outliers	Total n. of respondents	153
	<i>Outliers drop-out</i>	(1)
	<i>Protest responses drop-out</i>	(6)
SD Check	<i>Total n. of respondents</i>	146
	<i>Drop-out in this section</i>	(1)
Final Dataset	Final Valid Sample	145
	Human Condition	65
	AI Condition	80

4.4.3 Sample Characteristics

Following the data screening process detailed in the preceding sections, a comprehensive analysis of the final dataset (N = 145) was conducted to establish the demographic profile of the respondents. The final sample exhibits a relatively diverse distribution across five key demographic parameters: gender, age, nationality, education level, and occupational status (see Table 4.4).

Table 4.4. Sample Demographic Characteristics (N = 145) (Source: the author).

Category	Frequency (n)	%
Gender		
Male	47	32.41%
Female	93	64.14%
Prefer not to say	5	3.45%
Age Group		

18-24	38	26.21%
25-34	27	18.62%
35-44	25	17.24%
45-54	14	9.66%
55-64	27	18.62%
65+	14	9.66%
Nationality		
Italian	144	99.31%
Other	1	0.69%
Education Level		
Lower secondary education	19	13.10%
High school diploma	57	39.31%
Bachelor's degree	38	26.21%
Master's degree	28	19.31%
PhD	3	2.07%
Employment Status		
Student	31	21.38%
Employed (including self-employed)	85	58.62%
Unemployed	6	4.13%
Retired	18	12.41%
Other	5	3.45%

Overall, the sample is predominantly female, comprising 64.14% (n = 93) of the respondents, while male participants account for 32.41% (n = 47). In terms of age distribution, the sample is relatively well-distributed across all age groups, successfully capturing a wide range of adult consumers. While the largest single group consists of younger adults aged 18 to 24 (26.21%, n = 38), the remaining respondents are well-distributed across middle and late adulthood, with significant clusters in the 25–34

(18.62%, n = 27) and 55–64 (18.62%, n = 27) brackets. This distribution ensures that the study findings are not driven by a single age group, thereby enhancing the external validity of the results.

Aligning with the target population defined in Section 4.4.1, the sample is almost exclusively of Italian nationality (99.31%, n = 144). Furthermore, the respondents demonstrate a robust level of formal education and workforce participation. A substantial 47.59% (n = 69) of the sample has a university-level degree (Bachelor's, Master's, or PhD), while 39.31% (n = 57) have attained a high school diploma. From an occupational standpoint, the majority of the sample is actively engaged in the workforce, with 58.62% (n = 85) identifying as employed or self-employed, followed by students (21.38%, n = 31).

This specific demographic composition, characterized by educated, professionally active adult consumers, is well-suited for the context of this research. It represents a sample that possesses both the cognitive and emotional capabilities required to evaluate complex cues of brand authenticity, as well as the purchasing power necessary to formulate realistic economic evaluations regarding their willingness to pay for luxury products following exposure to an advertisement labelled as either traditional photography or AI-generated.

4.5 Data Analysis Strategy

The conceptual model developed in the present study will be empirically tested using Partial Least Squares Structural Equation Modelling (PLS-SEM), implemented via the SmartPLS 4 software. As briefly described in Section 4.1, the choice to use PLS-SEM over covariance-based SEM (CB-SEM) is justified by several factors. First, PLS-SEM is particularly suitable due to its inherent flexibility, its focus on prediction, and its suitability for exploratory research contexts. Because this study seeks to test and extend existing theory by applying the Stimulus-Organism-Response (S-O-R) framework to the novel domain of AI-generated advertising, while simultaneously explaining and

predicting key endogenous constructs, the variance-based PLS-SEM approach is deemed highly appropriate (Hair et al., 2017). Second, PLS-SEM is specifically designed to handle complex conceptual models involving a multitude of latent constructs and relationships, a capability that is essential for this study. Furthermore, PLS-SEM is advantageous because it is robust to non-normal data distributions, typical of Likert-scale survey data, which are treated as continuous variables within PLS-SEM (Hair et al., 2017).

Regarding the sample size adequacy, in the present model the maximum number of structural paths directed at a single endogenous construct is five. These five paths correspond to the two dependent variables (WTP and Brand Authenticity), each of which receives the direct effect of GenAI disclosure (H1a, H1b), the two mediating paths from Perceived Effort (H2b) and Moral Disgust (H3b), and the two control variables (Luxury Interest and Visual Appeal). According to the “10-times rule” proposed by Hair et al. (2017), the minimum sample size should be at least ten times the maximum number of structural paths directed at any single endogenous construct. Therefore, the present model requires a minimum sample size of 50 respondents. The final sample of 145 respondents exceeds this threshold, thereby ensuring robust statistical estimations.

As illustrated in Figure 3.1, the proposed conceptual model must be operationalized within the PLS-SEM framework. In PLS-SEM, variables are categorized as either exogenous or endogenous (Hair et al., 2017). The experimentally manipulated AI disclosure variable serves as the sole exogenous construct of the study. Conversely, the mediators (Perceived Effort and Moral Disgust) and the dependent variables (WTP and Brand Authenticity) operate as endogenous constructs as they are being explained and predicted by other variables in the model. Consumers' Need for Uniqueness acts as a moderator, while Visual Appeal and Luxury Interest are operationalized as control variables to isolate the true causal effects of the manipulation. All multi-item latent constructs within the structural model are operationalized as reflective measurement models. According to Hair et al. (2017), in reflective measurement models, observable indicators are assumed to be caused by the underlying construct, meaning that the latent variable drives the consumers' responses to the measurement items. There are two

exceptions to the multi-item reflective specification: the control measure for Luxury Interest and the dependent variable WTP. These variables are both operationalized as single-item observed variables. Since these variables are single-item measures, they are not subject to traditional internal consistency reliability and convergent validity assessments, and are instead directly included in the structural model (Hair et al., 2017).

Before analyzing the data, the survey data collected via Qualtrics were cleaned and prepared to ensure full compatibility with the SmartPLS software. Specifically, all measurement items were coded using acronyms (e.g., VIS_APP_1 for the first question of Visual Appeal). Furthermore, the AI-disclosure label variable was dummy coded (0 = Human, 1 = AI) to allow for proper categorical evaluation within the structural model. Table 4.5 provides a comprehensive overview of the constructs, items, and specific variable coding utilized for the PLS-SEM analysis.

Table 4.5. Overview of Constructs, Measurement Items, and Variable Coding for the PLS-SEM Analysis (Source: the author).

Constructs	Variable Coding	Measurement Items
AI Disclosure	[GROUP]	0: "Photography by G. Von Hassel"
		1: "Image entirely generated by Artificial Intelligence"
Luxury Interest	[LUX_INT]	How interested are you in the world of luxury products and accessories?
Visual Appeal	[VIS_APP]	
	[VIS_APP_1]	The way the product is displayed in the image is attractive.
	[VIS_APP_2]	The image is aesthetically appealing.
	[VIS_APP_3]	I like the way the image looks.
Perceived Effort	[PER_EFF]	
	[PER_EFF_1]	Time required to create the image.

	[PER_EFF_2]	Effort required to create the image.
	[PER_EFF_3]	Difficulty of creating the image.
Moral Disgust	[MOR_DISG]	
	[MOR_DISG_1]	To what extent does the way this image was created make you feel disgusted?
	[MOR_DISG_2]	To what extent does the way this image was created make you feel repulsed?
	[MOR_DISG_3]	To what extent does the way this image was created make you feel sickened?
Brand Authenticity	[BRAND_AUTH]	
	[BRAND_AUTH_1]	How authentic do you consider the brand to be?
	[BRAND_AUTH_2]	How trustworthy do you consider the brand to be?
	[BRAND_AUTH_3]	How credible do you consider the brand to be?
Willingness to Pay	[WTP]	What is the MAXIMUM price you would be willing to pay to purchase this bag?
Need for Uniqueness	[NEED_UNIQ]	
	[NEED_UNIQ_1]	I avoid products bought by the average consumer.
	[NEED_UNIQ_2]	I lose interest in brands that become too popular.
	[NEED_UNIQ_3]	The more common a product, the less I want it.

The data analysis presented in Chapter 5 will follow the two-step procedure recommended by Hair et al. (2017). First, the measurement model (outer model) will be evaluated to assess the reliability and validity of the constructs. This includes verifying indicator reliability through outer loadings, establishing internal consistency reliability via Cronbach's Alpha and Composite Reliability, confirming convergent validity using the Average Variance Extracted (AVE) and assessing discriminant validity using the Heterotrait-Monotrait (HTMT) ratio (Hair et al., 2017). Second, the structural model

(inner model) will be assessed to test the hypothesized relationships. This involves checking for collinearity using the Variance Inflation Factor (VIF), as well as evaluating the path coefficients (β) and their statistical significance using a nonparametric bootstrapping procedure. The bootstrapping procedure will be executed using 5,000 subsamples drawn with replacement from the original dataset of 145 valid observations, as established by Hair et al. (2017). Moreover, the model's predictive capabilities will be evaluated using the coefficient of determination (R^2), and effect sizes (f^2). Finally, the structural analysis will evaluate the specific moderation and parallel mediation effects outlined in the theoretical framework, with mediation type classified according to the framework proposed by Zhao et al. (2010).

The present chapter has established a rigorous experimental design and described the development of the advertising stimuli. It has also presented, defined, and validated the measurement scales and the corresponding questionnaire. Finally, it detailed the data collection and cleaning procedures, resulting in a robust and reliable final sample of 145 respondents. Chapter 5 will present the empirical findings of the PLS-SEM analysis, beginning with the measurement model assessment, and followed by the structural model evaluation and hypothesis testing.

5. Data Analysis and Results

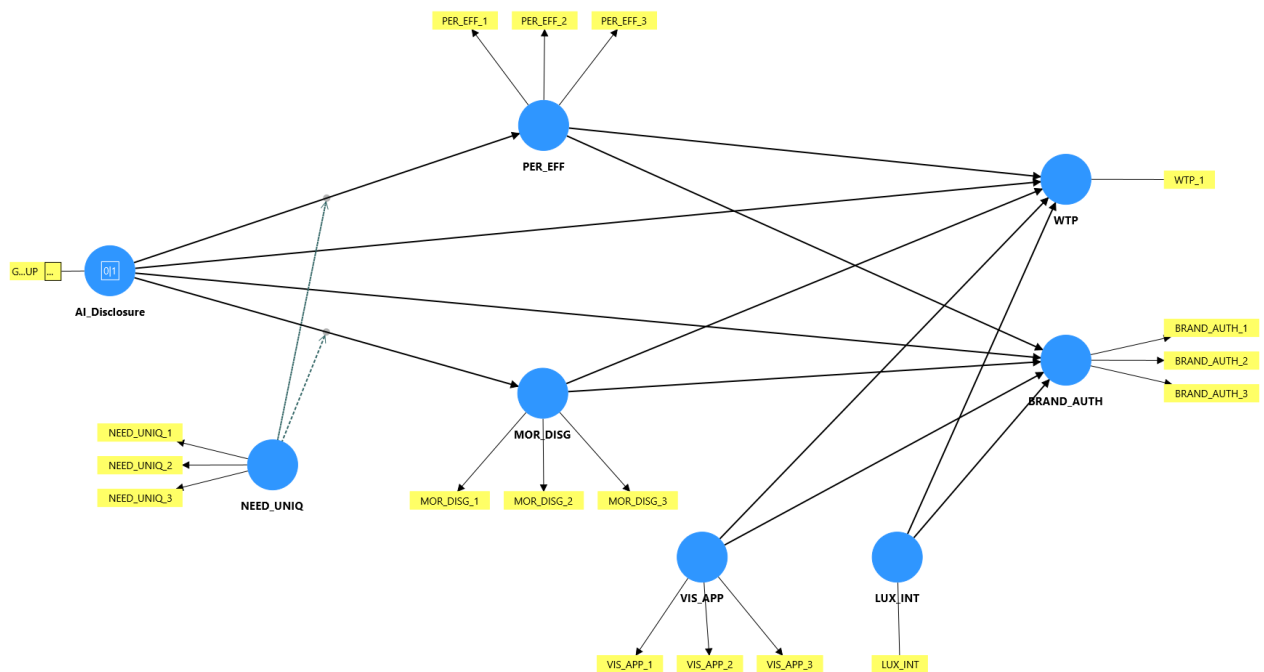
5.1 Introduction to the PLS-SEM Analysis

Building upon the methodological and analytical framework established in Chapter 4, the present chapter details the empirical estimation and evaluation of the proposed conceptual model. As established by Hair et al. (2017), a prerequisite of any Partial Least Squares Structural Equation Modelling (PLS-SEM) application is the correct specification of the measurement model, which necessitates a careful distinction between reflective and formative construct specifications. In reflective measurement models (Mode A), causality flows from the latent construct to its observable indicators. Therefore, the indicators act as interchangeable, highly covarying manifestations of a single underlying trait, meaning that removing one indicator does not alter the conceptual meaning of the construct (Hair et al., 2017). In formative measurement models (Mode B), the indicators are theorized to cause or form the construct, they do not need to be correlated, and each captures a conceptually unique dimension of the construct's domain (Hair et al., 2017).

In the present study, all multi-item latent constructs are specified as reflective measurement models. This is justified because their respective measurement items are co-varying manifestations of a single underlying psychological or perceptual state, rather than independent drivers of it. In the case of Moral Disgust, for instance, a consumer feeling disgusted, repulsed, and sickened represents three co-varying emotional expressions triggered by the same underlying affective state, rather than three independent causes combining to form it. Furthermore, the control variable Luxury Interest and Willingness to Pay (WTP) are operationalized as single-item observed variables. Similarly, the independent variable AI Disclosure is operationalized as a binary observed variable (dummy coded 0 = Human, 1 = AI-generated). These single-item and dummy-coded variables are not subject to reflective or formative measurement model specifications. Therefore, they are excluded from the outer model evaluation and are instead entered directly into the structural model, in accordance with the guidelines established by Hair et al. (2017).

The complete path model as implemented in SmartPLS 4 is depicted in Figure 5.1, illustrating the operationalization of the S-O-R framework developed in Chapter 3 as a structural model. Within the outer model, the five reflective constructs display outward-pointing arrows directed to their respective indicators, consistent with a Mode A reflective measurement model (Hair et al., 2017). The three directly observed variables, AI Disclosure, Luxury Interest, and WTP, enter the structural model without surrounding measurement models. The inner model reflects the dual-pathway mediation structure, positioning AI Disclosure as the sole exogenous stimulus, with Perceived Effort and Moral Disgust representing the two parallel mediating constructs and WTP and Brand Authenticity as the primary outcome variables. Need for Uniqueness is operationalized via interaction terms directed onto both S-O paths, capturing the moderated mediation. Furthermore, the control variables, Visual Appeal and Luxury Interest, are connected directly to both WTP and Brand Authenticity.

Figure 5.1. The Path Model created in SmartPLS 4 (Source: SmartPLS 4).



To execute the PLS-SEM analysis, the present chapter follows the two-step analytical procedure recommended by Hair et al. (2017). Specifically, this procedure requires that the measurement model (outer model) must be assessed first, followed by the evaluation of the structural model (inner model). The underlying logic is that the structural path coefficients are only meaningful and interpretable if the latent constructs they connect to have been empirically demonstrated to be both reliable and valid. A flawed measurement model introduces measurement error into the structural estimates. Therefore, successfully validating the outer model serves as a prerequisite before any empirical examination of the structural relationships can be undertaken. Because PLS-SEM makes no distributional assumptions and relies on a nonparametric approach, the statistical significance for all path coefficients, indirect effects, and interaction terms is assessed via a bootstrapping procedure rather than conventional parametric tests (Hair et al., 2017). Furthermore, the structural model assessed in Section 5.3 incorporates a dual-path parallel mediation structure and moderated mediation, requiring in particular the evaluation of specific indirect effects and the index of moderated mediation alongside the standard direct path coefficients (Hair et al., 2017; Hayes, 2015).

The remainder of this chapter is structured as follows. Section 5.2 evaluates the measurement model, verifying the reliability and validity of the reflective constructs. Specifically, this involves evaluating indicator reliability via outer loadings, internal consistency reliability via Cronbach's Alpha and Composite Reliability, convergent validity via the Average Variance Extracted (AVE), and discriminant validity via the Heterotrait-Monotrait (HTMT) ratio. Subsequently, Section 5.3 details the assessment of the structural model, evaluating collinearity via the Variance Inflation Factor (VIF), the significance and relevance of path coefficients via bootstrapping, explanatory power via the Coefficient of Determination (R^2), and Effect Sizes via f^2 . Additionally, Section 5.3 addresses the formal testing of the direct hypotheses, alongside the parallel mediation and moderated mediation analyses. Section 5.4 discusses the WTP results distribution and, lastly, Section 5.5 provides a consolidated summary of the hypothesis testing results.

5.2 Assessment of the Measurement Model (Outer Model)

The analysis of the path model must begin with the evaluation of the outer measurement model. This initial step is necessary to establish two foundational properties of the measurement model: reliability and validity. In PLS-SEM, reliability indicates the consistency and stability with which a set of indicators measures its intended latent construct, ensuring the measurement is free from random error. Conversely, validity indicates how accurately those indicators capture the theoretical construct (Hair et al., 2017). Validity comprises convergent validity (meaning indicators of the same construct share sufficient variance) and discriminant validity (which ensures the constructs are empirically distinct from one another) (Hair et al., 2017). A measurement model that lacks these properties risks producing relationships in the structural model that reflect measurement artefacts. Therefore, ensuring the absence of substantial measurement error is a prerequisite (Hair et al., 2017).

The PLS-SEM algorithm was run using the SmartPLS 4 software (Hair et al., 2022). The initial estimation of the measurement model relied on the path weighting scheme, which is the default and recommended method for reflective PLS-SEM path model estimation. This scheme accounts for directionality of the structural model relationships when computing construct scores to produce superior predictive accuracy (Hair et al., 2017). All data were standardized by SmartPLS 4 prior to estimation, ensuring that the main effects of the predictor and moderator are interpretable as simple effects, evaluated at the mean value of the moderator, rather than at an arbitrary zero point (Hair et al., 2017). Moreover, the algorithm was configured with a maximum of 300 iterations and a stop criterion of 10^{-7} and converged after only two iterations, indicating a well-specified and stable model structure (Hair et al., 2017). No missing values were present in the dataset, as the final sample of 145 observations had already been subject to pre-processing during the data preparation stage detailed in Section 4.4.2. Accordingly, missing values in the SmartPLS 4 settings were designated with a placeholder value of -99, consistent with standard conventions (Hair et al., 2022).

Running the algorithm provided the outer loadings for each reflective construct, serving as the primary input for the measurement model assessment. In addition, a bootstrapping procedure with 5,000 subsamples and a significance level of 0.05 was performed at this stage with a fixed seed. This bootstrapping procedure was employed exclusively to derive Bias-Corrected and Accelerated (BCa) bootstrap confidence intervals for the Heterotrait-Monotrait (HTMT) discriminant validity criterion, in accordance with the configuration detailed in Section 5.2.4 (Hair et al., 2017; Hair et al., 2022).

Accordingly, the measurement model assessment is structured across four sequential sub-sections, each addressing a distinct dimension of construct quality. The evaluation criteria are applied exclusively to the five reflective multi-item constructs. The three directly observed variables, AI Disclosure, Luxury Interest, and WTP, are not subject to outer model assessment and are therefore excluded from the analyses presented in this section. The evaluation proceeds sequentially, beginning with indicator reliability (Section 5.2.1), followed by internal consistency reliability (Section 5.2.2), convergent validity (Section 5.2.3), and discriminant validity (Section 5.2.4). A comprehensive summary of the measurement model assessment is provided in Section 5.2.5.

5.2.1 Indicator Reliability (Outer Loadings)

Indicator reliability within reflective PLS-SEM measurement models is assessed through outer loadings, which represent the standardized regression coefficients linking each indicator to its underlying latent construct (Hair et al., 2017). In reflective measurement models, each outer loading captures the strength of the relationship between an indicator and its construct (Hair et al., 2017). The square of a standardized outer loading is referred to as the communality of the indicator and represents the proportion of the indicator's variance that is explained by the latent construct (Hair et al., 2017). Specifically, an outer loading of 0.708 implies a communality of approximately 0.50, meaning the construct accounts for at least half of the indicator's variance, with the remaining variance attributable to measurement error. The established rule of thumb requires that all

standardized outer loadings meet or exceed the threshold of 0.708, as values below this level indicate that measurement error accounts for the majority of the indicator's variance (Hair et al., 2017). SmartPLS 4 produces outer loading estimates for all constructs and observed variables in the path model, including single-item observed variables for which the loading is fixed at 1.00 by definition and carries no interpretive value for reliability assessment. The present evaluation is made only on the outer loadings of the five reflective multi-item constructs (Hair et al., 2017).

Hair et al. (2017) establish additional considerations for the treatment of indicators whose loadings fall below the primary threshold of 0.708. They recommend retaining indicators that fall between 0.40 and 0.70 only when their removal does not improve composite reliability or AVE (assessed in Section 5.2.3) above the established thresholds, and when content validity considerations support their inclusion. However, indicators with outer loadings below 0.40 should always be removed (Hair et al., 2017). Each indicator's outer loading on its assigned construct is expected to exceed its correlations with all other constructs in the model. This is commonly referred to as cross-loading examination and is formally assessed within the discriminant validity evaluation presented in Section 5.2.4 (Hair et al., 2017). The outer loadings of the fifteen indicators across the five reflective constructs are reported individually in Table 5.1.

As reported in Table 5.1, all fifteen indicators exceed the 0.708 threshold, demonstrating satisfactory indicator reliability (Hair et al., 2017). The outer loadings across all fifteen indicators range from a minimum of 0.769 (MOR_DISG_1) to a maximum of 0.945 (BRAND_AUTH_2). Within this distribution, two specific measurement patterns warrant acknowledgement.

First, MOR_DISG_1, at 0.769, is lower than the other two indicators of the same construct (MOR_DISG_2 = 0.925; MOR_DISG_3 = 0.821), indicating meaningful heterogeneity in the loading structure of Moral Disgust. Nevertheless, MOR_DISG_1 exceeds the 0.708 threshold, and no further action is required (Hair et al., 2017). Similarly, the Need for Uniqueness construct exhibits meaningful heterogeneity across its indicator loadings (NEED_UNIQ_1 = 0.859; NEED_UNIQ_2 = 0.784; NEED_UNIQ_3 = 0.831). Nevertheless,

NEED_UNIQ_2 similarly exceeds the 0.708 threshold and no remedial action is required (Hair et al., 2017). The loading heterogeneity observed within both constructs has implications for the internal consistency reliability estimates and is examined in Section 5.2.2.

Table 5.1. *Outer loadings of the indicators (Source: SmartPLS 4).*

Indicator	BRAND_AUTH	MOR_DISG	NEED_UNIQ	PER_EFF	VIS_APP
BRAND_AUTH_1	0.878				
BRAND_AUTH_2	0.945				
BRAND_AUTH_3	0.944				
MOR_DISG_1		0.769			
MOR_DISG_2		0.925			
MOR_DISG_3		0.821			
NEED_UNIQ_1			0.859		
NEED_UNIQ_2			0.784		
NEED_UNIQ_3			0.831		
PER_EFF_1				0.932	
PER_EFF_2				0.919	
PER_EFF_3				0.865	
VIS_APP_1					0.900
VIS_APP_2					0.911
VIS_APP_3					0.909

Second, BRAND_AUTH_2 (0.945) and BRAND_AUTH_3 (0.944) show near-identical and high outer loadings. While both values satisfy the indicator reliability criterion, their near-identical loadings suggest that these two indicators may be capturing a very narrow and overlapping facet of the Brand Authenticity construct (Hair et al., 2017). The construct-

level implications of this loading convergence are examined in the internal consistency reliability assessment in Section 5.2.2. Accordingly, the indicator reliability criterion is satisfied for all fifteen indicators. Grounded in these results, the analysis proceeds to the assessment of internal consistency reliability in Section 5.2.2.

5.2.2 Internal consistency (Cronbach's Alpha, Composite Reliability)

Internal consistency reliability in PLS-SEM is traditionally assessed through Cronbach's alpha (Hair et al., 2017). This metric estimates reliability based on the intercorrelations among the observed indicator variables of each reflective construct, operating under the assumption that all indicators are equally reliable and possess equal outer loadings. This equal-weighting assumption, however, is routinely unsatisfied in PLS-SEM, where indicators characteristically exhibit heterogeneous outer loadings. Consequently, Cronbach's alpha tends to underestimate true internal consistency reliability and is best regarded as a conservative lower bound (Hair et al., 2017). Cronbach's alpha values between 0.70 and 0.90 are generally considered satisfactory in confirmatory and explanatory research contexts, whereas values between 0.60 and 0.70 are often deemed acceptable in exploratory contexts. Conversely, values exceeding 0.95 are typically considered undesirable since they suggest semantic redundancy among the indicators, implying that the items may simply be capturing overlapping aspects of the same underlying phenomenon (Hair et al., 2017).

To address the limitations of Cronbach's alpha within the PLS-SEM context, composite reliability (ρ_c) explicitly accounts for the differing outer loadings of each indicator and produces a more accurate estimate of internal consistency (Hair et al., 2017). This measure is recognized as the methodologically preferred measure of internal consistency reliability in PLS-SEM. The same evaluative thresholds apply to composite reliability, with values between 0.70 and 0.90 generally considered satisfactory and values above 0.95 deemed undesirable (Hair et al., 2017). In addition to these two metrics, the reliability coefficient (ρ_A) serves as a third complementary measure designed to address the respective limitations of both Cronbach's alpha and composite reliability (Hair et al., 2017; Hair et al., 2022). These three measures establish a plausible interval within which

true reliability is likely to fall: Cronbach's alpha represents the lower bound of true reliability, ρ_c represents the upper bound, and ρ_A , which is more sensitive to heterogeneity in indicator loadings, is positioned between the two, providing a precise estimate of true reliability (Hair et al., 2017). Therefore, reporting all three measures is considered a good practice in PLS-SEM analysis, as they collectively define the plausible range within which true reliability lies. Grounded in this methodological framework, the internal consistency reliability of the five reflective constructs is assessed and reported in Table 5.2.

Table 5.2. Cronbach's alpha, Composite Reliability, and AVE of the Measurement Model (Source: SmartPLS 4).

Construct	Cronbach's alpha	Reliability coefficient (ρ_A)	Composite reliability (ρ_c)	Average variance extracted (AVE)
BRAND_AUTH	0.913	0.930	0.945	0.852
MOR_DISG	0.796	0.898	0.878	0.707
NEED_UNIQ	0.791	0.880	0.865	0.681
PER_EFF	0.890	0.895	0.932	0.820
VIS_APP	0.892	0.893	0.933	0.822

As reported in Table 5.2, all five reflective constructs satisfy the established reliability thresholds across the three measures: Cronbach's alpha, ρ_A , and composite reliability (ρ_c). Specifically, the reliability estimates fall within or marginally above the satisfactory range of 0.70 to 0.90, with no construct exceeding the 0.95 ceiling that would indicate undesirable semantic redundancy among the indicators (Hair et al., 2017). Furthermore, while the Average Variance Extracted (AVE) column is included in Table 5.2 for completeness, the interpretation of these values is assessed in Section 5.2.3.

Within these satisfactory results, two methodological nuances warrant acknowledgement, neither of which suggests a validity failure. For two constructs, Moral Disgust (MOR_DISG: $\rho_A = 0.898 > \rho_c = 0.878$) and Need for Uniqueness (NEED_UNIQ: $\rho_A = 0.880 > \rho_c = 0.865$), ρ_A exceeds ρ_c , a pattern that is theoretically atypical given that ρ_c is conventionally positioned as the upper bound of true reliability. However, this pattern

arises because ρ_A is sensitive to heterogeneity in outer loadings. When outer loadings within a construct differ substantially, ρ_A can inflate estimates beyond ρ_C (Hair et al., 2017; Hair et al., 2022). As established in Section 5.2.1, both constructs exhibit meaningful heterogeneity across their indicator loadings. Nonetheless, this pattern does not constitute a validity problem, as all reliability values for both constructs remain within the acceptable 0.70 to 0.95 range (Hair et al., 2017). Furthermore, Brand Authenticity's composite reliability of 0.945 closely approaches the 0.95 upper threshold, though it remains within the acceptable range (Hair et al., 2017). The results of Table 5.2 confirm that all five reflective constructs satisfy the internal consistency reliability criteria. Therefore, the analysis proceeds with the assessment of convergent validity in Section 5.2.3.

5.2.3 Convergent Validity (Average Variance Extracted)

Convergent validity represents the extent to which the indicators of a reflective construct share sufficient common variance, thereby confirming that they are measuring the same underlying latent concept (Hair et al., 2017). In PLS-SEM, convergent validity is assessed with the AVE. Single-item observed variables (for which the indicator loading is fixed at 1.00) yield an AVE of 1.00 and are therefore excluded, restricting this assessment exclusively to the five reflective multi-item constructs (Hair et al., 2017). While communality was assessed at the individual indicator level in Section 5.2.1, AVE extends this logic to the construct level by computing the mean of the squared outer loadings across all indicators of a given construct (Hair et al., 2017). An AVE of 0.50 or higher indicates that the construct explains more variance in its indicators than is attributable to measurement error, and is accordingly considered the minimum threshold for acceptable convergent validity (Hair et al., 2017). AVE constitutes a more stringent criterion than composite reliability because it relies on the mean of squared loadings rather than a weighted sum; consequently, a construct may satisfy the internal consistency reliability thresholds established in Section 5.2.2 while failing the AVE criterion (Hair et al., 2017).

As reported in Table 5.2, all five constructs satisfy the AVE threshold of 0.50, ranging from a minimum of 0.681 (NEED_UNIQ) to a maximum of 0.852 (BRAND_AUTH). Notably,

BRAND_AUTH (0.852), PER_EFF (0.820), and VIS_APP (0.822) demonstrate particularly strong convergent validity with AVE values exceeding 0.80, indicating that each construct explains a substantial majority of its indicators' variance (Hair et al., 2017). Conversely, MOR_DISG (0.707) and NEED_UNIQ (0.681) record lower but acceptable AVE values. Specifically, NEED_UNIQ's lower AVE is partially driven by the loading heterogeneity already established in Section 5.2.1, primarily reflecting NEED_UNIQ_2's weaker outer loading of 0.784 (Hair et al., 2017). Nonetheless, an AVE value of 0.681 exceeds the established threshold by a meaningful margin and does not constitute a convergent validity concern. Taken together, convergent validity is established for all reflective constructs, and the analysis proceeds to the discriminant validity assessment.

5.2.4 Discriminant Validity (Cross-loadings, Fornell-Larcker, HTMT)

Discriminant validity represents the extent to which a construct is empirically distinct from all other constructs in the model, confirming that it captures phenomena not represented by any other latent variable (Hair et al., 2017). In the present study, discriminant validity is evaluated through three sequential criteria: Cross-Loadings, the Fornell-Larcker criterion, and the Heterotrait-Monotrait (HTMT) ratio. These criteria are assessed in order of increasing methodological robustness, in accordance with Hair et al. (2017). Therefore, the assessment of discriminant validity begins by analyzing the cross-loadings.

Cross-loadings represent an indicator-level criterion for discriminant validity. Specifically, each indicator is expected to load most strongly on its assigned construct and to exhibit substantially lower correlations (referred to as cross-loadings) with all other constructs in the model (Hair et al., 2017). An indicator's cross-loading on another construct that exceeds or approaches its loading on the assigned construct constitutes a violation of this pattern, indicating discriminant validity failure at the indicator level (Hair et al., 2017). Specifically, each indicator's outer loading on its assigned construct must exceed all of its cross-loadings with every other construct in the model. However, Hair et al. (2017) note that this criterion lacks the sensitivity to detect discriminant validity problems when two constructs are highly correlated, limiting cross-loadings' reliability

as a standalone assessment. Therefore, this criterion serves as an initial check rather than a definitive test of discriminant validity. This limitation motivates the subsequent application of the Fornell-Larcker criterion and the HTMT ratio in the assessments that follow. The cross-loading matrix for the fifteen indicators across the five reflective constructs is reported in Table 5.3.

Table 5.3. Cross-loadings of the Measurement Model (Source: SmartPLS 4).

Indicator	BRAND_AUTH	MOR_DISG	NEED_UNIQ	PER_EFF	VIS_APP
BRAND_AUTH_1	0.878	-0.223	-0.002	0.264	0.245
BRAND_AUTH_2	0.945	-0.268	-0.074	0.276	0.326
BRAND_AUTH_3	0.944	-0.265	-0.039	0.338	0.300
MOR_DISG_1	-0.204	0.769	-0.121	0.069	-0.188
MOR_DISG_2	-0.304	0.925	-0.119	-0.123	-0.356
MOR_DISG_3	-0.146	0.821	-0.184	0.110	-0.270
NEED_UNIQ_1	-0.016	-0.170	0.859	0.044	-0.042
NEED_UNIQ_2	-0.106	-0.074	0.784	-0.080	-0.084
NEED_UNIQ_3	-0.026	-0.115	0.831	0.014	-0.001
PER_EFF_1	0.242	0.028	-0.039	0.932	0.056
PER_EFF_2	0.319	-0.060	0.008	0.919	0.124
PER_EFF_3	0.312	0.001	0.071	0.865	0.198
VIS_APP_1	0.288	-0.357	0.009	0.180	0.900
VIS_APP_2	0.297	-0.258	-0.101	0.149	0.911
VIS_APP_3	0.277	-0.299	-0.025	0.037	0.909

As reported in Table 5.3, all fifteen indicators load most strongly on their designated constructs, with primary loadings ranging from 0.769 to 0.945 and no cross-loading exceeding 0.356 in absolute terms (Hair et al., 2017). For instance, the VIS_APP indicators'

primary loadings on Visual Appeal, ranging from 0.900 to 0.911, substantially exceed their highest cross-loadings on any other construct, with the largest cross-loading recorded at 0.297 on Brand Authenticity. Similarly, the PER_EFF indicators' primary loadings on Perceived Effort, ranging from 0.865 to 0.932, far exceed their highest cross-loadings on any other construct in the model. These results provide an initial indicative support for discriminant validity at the indicator level. Nevertheless, given the acknowledged sensitivity limitations of cross-loadings in detecting discriminant validity problems when constructs are moderately correlated, the assessment of discriminant validity proceeds to the Fornell-Larcker criterion (Hair et al., 2017).

The Fornell-Larcker criterion represents a more rigorous assessment of discriminant validity than cross-loadings, and it operates at a construct level rather than at an indicator level (Hair et al., 2017). For this criterion, a construct should share more variance with its own indicators than it shares with any other construct in the model. The criterion does so by comparing each construct's internal shared variance, captured by the AVE, against its squared correlations with all other constructs (Hair et al., 2017). Specifically, discriminant validity is established when the square root of each construct's AVE exceeds its correlation with every other construct in the model (Hair et al., 2017). The comparison is presented in a matrix format in which the square roots of the AVE values appear on the diagonal and the inter-construct correlations appear in the off-diagonal positions. Discriminant validity is supported when each diagonal value exceeds all values in its corresponding row and column (Hair et al., 2017). The Fornell-Larcker results for the five reflective constructs are reported in Table 5.4 against which the criterion is now assessed.

As reported in Table 5.4, all five diagonal values exceed all off-diagonal inter-construct correlations in their respective rows and columns, satisfying the Fornell-Larcker criterion across the full matrix (Hair et al., 2017). The diagonal values range from 0.825 (NEED_UNIQ) to 0.923 (BRAND_AUTH), and the highest absolute off-diagonal correlation is 0.336 (between VIS_APP and MOR_DISG). Although several off-diagonal correlations carry negative signs, these are theoretically consistent with the directional relationships

proposed in the S-O-R framework and do not affect the criterion's evaluation, which is assessed in absolute value terms.

Table 5.4. *Fornell-Larcker of the Measurement Model (Source: SmartPLS 4).*

Construct	BRAND_AUTH	MOR_DISG	NEED_UNIQ	PER_EFF	VIS_APP
BRAND_AUTH	0.923				
MOR_DISG	-0.275	0.841			
NEED_UNIQ	-0.044	-0.160	0.825		
PER_EFF	0.319	-0.010	0.013	0.906	
VIS_APP	0.317	-0.336	-0.044	0.136	0.907

Therefore, the Fornell-Larcker criterion is satisfied for all five reflective constructs, providing construct-level empirical support for discriminant validity (Hair et al., 2017). Nevertheless, the Fornell-Larcker criterion is known to perform poorly when indicator loadings across constructs are of broadly similar magnitude (Hair et al., 2017; Henseler et al., 2015). This condition characterizes the present measurement model where outer loadings range from 0.769 to 0.945. This is a range sufficiently homogeneous to reduce the criterion's sensitivity to potential construct overlap (Hair et al., 2017; Henseler et al., 2015). Consequently, the HTMT ratio is applied as the third and primary criterion for discriminant validity assessment, having demonstrated greater statistical sensitivity in detecting construct overlap relative to both cross-loadings and the Fornell-Larcker criterion (Hair et al., 2017; Henseler et al., 2015).

The Heterotrait-Monotrait (HTMT) ratio, proposed by Henseler et al. (2015) as a superior alternative to cross-loadings and the Fornell-Larcker criterion, is the primary and most methodologically rigorous criterion for discriminant validity assessment in PLS-SEM. Specifically, it provides an estimate of what the true correlation between two constructs would be if they were measured without error, referred to as the disattenuated

correlation, thereby offering a more precise test of construct distinctiveness (Hair et al., 2017; Henseler et al., 2015). HTMT is defined as the ratio of the average between-construct indicator correlations to the average within-construct indicator correlations, such that values approaching 1.00 indicate that two constructs are empirically indistinguishable, while substantially lower values confirm their discriminant validity (Henseler et al., 2015). HTMT values below 0.85 are generally required when the constructs assessed are conceptually distinct, while a less conservative threshold of 0.90 is applied when constructs are conceptually similar (Hair et al., 2017; Henseler et al., 2015). In the present study, all constructs are theoretically distinct and, consequently, the more conservative threshold of 0.85 applies, with values at or above this threshold indicating insufficient discriminant validity. The HTMT point estimate values for all construct pairs in the present study are reported in Table 5.5.

Table 5.5. HTMT of the Measurement Model (Source: SmartPLS 4).

Construct	BRAND_AUTH	MOR_DISG	NEED_UNIQ	PER_EFF	VIS_APP
BRAND_AUTH					
MOR_DISG	0.302				
NEED_UNIQ	0.082	0.192			
PER_EFF	0.354	0.142	0.086		
VIS_APP	0.348	0.383	0.076	0.162	

As reported in Table 5.5, all ten HTMT values fall well below the 0.85 threshold, which is generally considered to provide strong point-estimate support for discriminant validity across the full construct matrix (Hair et al., 2017; Henseler et al., 2015). The values range from a minimum of 0.076 (VIS_APP / NEED_UNIQ) to a maximum of 0.383 (VIS_APP / MOR_DISG). Consequently, the substantial margin between the highest HTMT value and the threshold of 0.85 indicates that no construct pair in the model approaches empirical indistinguishability, confirming that each reflective construct captures a theoretically distinct dimension within the S-O-R framework (Henseler et al., 2015). Furthermore,

while point estimates provide initial support for discriminant validity, Hair et al. (2017) recommend complementing this assessment with bootstrapped confidence intervals as a more precise inferential test.

Since PLS-SEM makes no distributional assumptions and is inherently nonparametric, standard parametric significance tests cannot be applied to determine whether HTMT values are significantly different from 1.00 (Hair et al., 2017; Henseler et al., 2015). Therefore, a bootstrapping procedure is typically employed to derive a sampling distribution of the HTMT statistic from which confidence intervals can be constructed to test construct distinctiveness inferentially (Hair et al., 2017). The SmartPLS 4 bootstrapping procedure in this study is configured in accordance with the recommendations of Hair et al. (2017) and Hair et al. (2022), drawing 5,000 subsamples, utilizing the complete bootstrapping option with no sign changes, and estimating Bias-Corrected and Accelerated (BCa) bootstrap confidence intervals based on two-tailed testing at a significance level of 0.05, with a fixed random seed for reproducibility. Discriminant validity between two constructs is confirmed when the upper bound of the bootstrapped confidence interval of their HTMT statistic falls below 1.00; conversely, a confidence interval that includes 1.00 indicates that the two constructs cannot be considered empirically distinct (Hair et al., 2017; Henseler et al., 2015). The bootstrapped BCa confidence intervals for all ten construct pairs are reported in Table 5.6.

As reported in Table 5.6, for all ten reflective construct pairs, the upper bound of the BCa bootstrapped confidence interval falls below 1.00, with upper bounds ranging from 0.093 (VIS_APP / NEED_UNIQ) to 0.571 (VIS_APP / MOR_DISG), thereby satisfying the inferential discriminant validity criterion for all pairs of constructs (Hair et al., 2017; Henseler et al., 2015). The highest upper bound is observed for the VIS_APP/MOR_DISG pair (97.5% = 0.571), which, despite representing the most extreme case, remains well below the value of 1.00. Consequently, the bootstrapped confidence interval assessment provides strong inferential support for discriminant validity across all five reflective constructs, confirming that each construct captures a theoretically distinct dimension of the conceptual model (Hair et al., 2017; Henseler et al., 2015).

Table 5.6. Bias-corrected and accelerated bootstrap confidence intervals of the Measurement Model
(Source: SmartPLS 4).

Construct pair	Original sample (O)	Sample mean (M)	Bias	2.5%	97.5%
MOR_DISG <-> BRAND_AUTH	0.302	0.311	0.009	0.152	0.426
NEED_UNIQ <-> BRAND_AUTH	0.082	0.134	0.053	0.026	0.119
NEED_UNIQ <-> MOR_DISG	0.192	0.214	0.022	0.077	0.346
PER_EFF <-> BRAND_AUTH	0.354	0.353	-0.001	0.176	0.519
PER_EFF <-> MOR_DISG	0.142	0.172	0.030	0.065	0.205
PER_EFF <-> NEED_UNIQ	0.086	0.140	0.054	0.036	0.106
VIS_APP <-> BRAND_AUTH	0.348	0.348	0.000	0.152	0.512
VIS_APP <-> MOR_DISG	0.383	0.382	-0.001	0.165	0.571
VIS_APP <-> NEED_UNIQ	0.076	0.141	0.065	0.026	0.093
VIS_APP <-> PER_EFF	0.162	0.181	0.019	0.072	0.296

Taken together, these results confirm that all four measurement model criteria, i.e., indicator reliability, internal consistency reliability, convergent validity, and discriminant validity, are satisfactorily met across all five reflective constructs. A comprehensive summary of this assessment is provided in Section 5.2.5.

5.2.5 Summary of the Measurement Model Assessment

The assessment of indicator reliability and internal consistency reliability across the five reflective constructs, Visual Appeal, Perceived Effort, Moral Disgust, Brand Authenticity, and Need for Uniqueness, demonstrated satisfactory results against the established criteria (Hair et al., 2017). All fifteen indicators across the five constructs exhibit outer loadings that meet or exceed the established threshold, confirming that each item reliably reflects its assigned latent construct, as reported in Table 5.1 (Hair et al., 2017). Furthermore, Cronbach's alpha, ρ_A , and ρ_C values satisfy the recommended range across

all constructs, as reported in Section 5.2.2. Two measurement nuances warrant acknowledgement, neither of which constitutes a validity concern: the ρ_A exceeding ρ_C for two constructs and BRAND_AUTH's composite reliability approaching the upper threshold. Moreover, AVE values for all five constructs exceed the 0.50 threshold, confirming that each construct explains more variance in its indicators than is attributable to measurement error, as established in Section 5.2.3 (Hair et al., 2017). Building upon these findings, the discriminant validity assessment conducted across three criteria (Cross-Loadings, the Fornell-Larcker criterion, and the HTMT ratio) confirms that all five reflective constructs are empirically distinct from one another, as reported in Section 5.2.4. Furthermore, the HTMT bootstrapped confidence interval assessment provides definitive inferential confirmation that each construct is empirically distinct from all others in the model (Hair et al., 2017; Henseler et al., 2015).

Notably, no indicator or construct was removed from the measurement model during the evaluation. Therefore, the original measurement specification proposed in Chapter 4 is fully retained. With the reliability and validity of the measurement model conclusively established across all four assessment criteria, the analysis proceeds to the evaluation of the structural model and formal hypothesis testing in Section 5.3.

5.3 Assessment of the Structural Model (Inner Model)

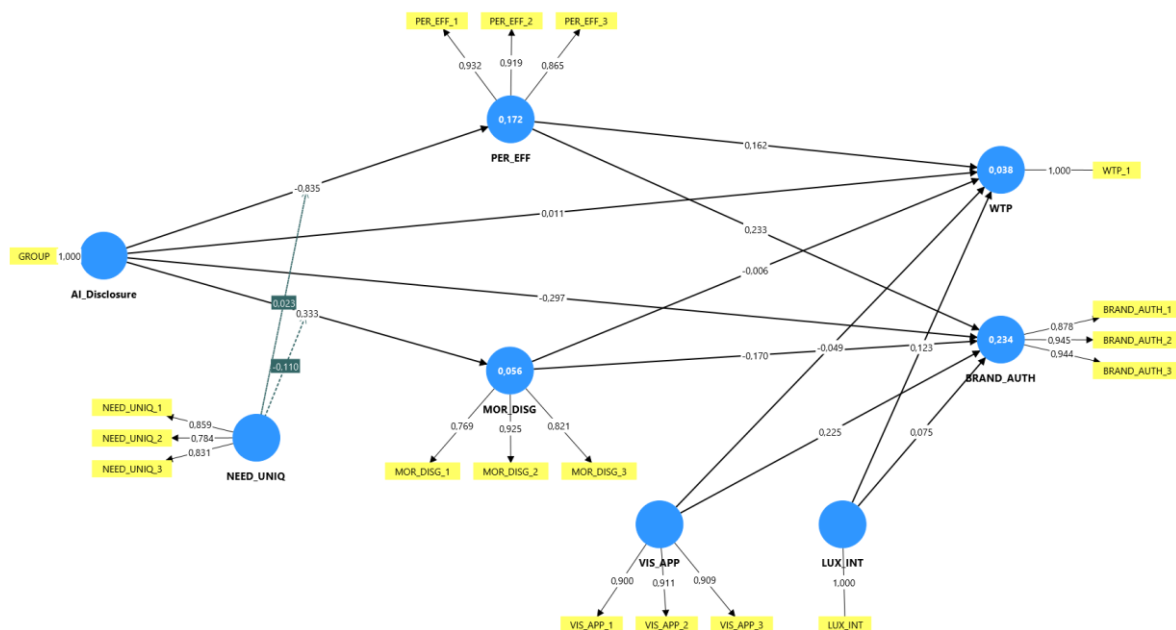
The completion of the measurement model fulfills the prerequisite for proceeding to the structural model evaluation, as path coefficients are interpretable only when the underlying constructs have been demonstrated to be reliable and valid (Hair et al., 2017). Accordingly, the structural model is defined as the inner component of the PLS path model that encodes the theorized cause-and-effect relationships among the latent constructs, and whose evaluation determines whether the hypothesized theoretical framework receives empirical support (Hair et al., 2017). The structural model of the present study is grounded in the S-O-R framework, assessing consumer reactions across the two experimental conditions. The structural evaluation encompasses not only direct path

coefficients but also the specific indirect effects via the two parallel mediating constructs and the index of moderated mediation. Furthermore, as established in Section 5.1, AI_Disclosure enters the model as a dummy-coded binary variable (0 = human, 1 = AI) without a surrounding reflective measurement model. Consequently, its path coefficients reflect the mean difference in the relevant endogenous construct between the two experimental conditions, evaluated at the mean value of the moderator following standardization (Hair et al., 2017). To evaluate these structural relationships, the analysis follows the structural model evaluation procedure established by Hair et al. (2017). The assessment begins with the examination of potential collinearity issues among the predictors employing the Variance Inflation Factor (VIF), as path coefficients in PLS-SEM are estimated through a series of ordinary least squares (OLS) regressions (Hair et al., 2017). Subsequently, the significance and relevance of the path coefficients are evaluated through a nonparametric bootstrapping procedure utilizing 5,000 subsamples, as established in Section 5.1 (Hair et al., 2017). Following the hypothesis testing, the structural model's in-sample explanatory power is assessed by examining the coefficient of determination (R^2) and the adjusted R^2_{adj} for the endogenous constructs. Lastly, the evaluation incorporates f^2 effect sizes to quantify the substantive contribution of each exogenous construct to the R^2 values of its corresponding endogenous targets (Hair et al., 2017).

While the methodology originally planned to include the Stone-Geisser Q^2 and q^2 statistics to assess predictive relevance, Hair et al. (2022) explicitly advise against the use of Stone-Geisser Q^2 , noting that it conflates in-sample and out-of-sample predictive power, and recommend the PLSpredict procedure instead. Nevertheless, applying PLSpredict requires k-fold holdout partitioning, which is methodologically impractical given the relatively small sample size of $N = 145$ and the complexity of the specified model. Furthermore, as detailed in Chapter 4, the application of PLS-SEM in the present study is driven by its capacity to handle complex structural models and to maximize the explained variance in the endogenous constructs, rather than by a commitment to out-of-sample predictive assessment (Hair et al., 2022). Consequently, the R^2 and f^2 metrics are

sufficient for the study's evaluative purposes, providing an adequate assessment of the model's explanatory capacity (Hair et al., 2017; Hair et al., 2022). Figure 5.2 presents the path model previously introduced in Figure 5.1 and displays the outer loadings on the indicator arrows, the path coefficients on the structural arrows, and the R^2 values within the endogenous construct circles (Hair et al., 2017). The specific values depicted in Figure 5.2 are examined in detail in further sections, within each corresponding criterion. Accordingly, the structural model assessment is structured as follows. The evaluation begins with the Collinearity Assessment (Section 5.3.1), followed by the analysis of Path Coefficients and Direct Effects (Section 5.3.2) and Explanatory Power (Section 5.3.3). Subsequently, the analysis addresses the parallel mediation of Perceived Effort and Moral Disgust (Section 5.3.4) and the moderated mediation of Need for Uniqueness (Section 5.3.5), before concluding with the assessment of the control variables, Luxury Interest and Visual Appeal (Section 5.3.6).

Figure 5.2. The Path Model created in SmartPLS 4 with the outer loadings (Source: SmartPLS 4).



5.3.1 Collinearity Assessment (Variance Inflation Factor)

In accordance with Hair et al. (2017), before examining path coefficients, the structural model must be assessed for potential collinearity among the predictors. Collinearity arises when two or more predictor constructs within the same regression equation are highly intercorrelated and may cause biased and unstable parameter estimates (Hair et al., 2017). In PLS-SEM, the estimation of path coefficients relies on OLS regressions, where each endogenous construct is regressed on its corresponding predictor constructs. To assess the presence of collinearity in a PLS-SEM structural model, the VIF is employed to measure collinearity among the structural model predictors. VIF quantifies the extent to which the variance of an estimated path coefficient is inflated due to its correlation with other predictors in the same equation (Hair et al., 2017). Higher VIF values generally indicate greater collinearity and a greater risk of biased structural estimates. Hair et al. (2017) establish a VIF of less than 5 as the primary threshold below which collinearity is not generally considered a critical concern, while Hair et al. (2022) propose the more conservative threshold of a VIF below 3.3, which is used for confirmatory and explanatory research contexts. Accordingly, the present study adopts a VIF of less than 5 as the main criterion, treating the VIF below 3.3 as a supplementary reference point (Hair et al., 2017; Hair et al., 2022). The VIF values for all predictor constructs across each set of structural model equations are reported in Table 5.7.

Table 5.7. Inner VIF Values of the Structural Model (Source: SmartPLS 4).

Constructs	BRAND_AUTH	MOR_DISG	PER_EFF	WTP
AI_Disclosure	1.272	1.009	1.009	1.272
LUX_INT	1.036			1.036
MOR_DISG	1.176			1.176
NEED_UNIQ		2.107	2.107	
PER_EFF	1.241			1.241
VIS_APP	1.161			1.161
NEED_UNIQ x AI_Disclosure		2.096	2.096	

As reported in Table 5.7, all VIF values fall below both the threshold of 5.0 established by Hair et al. (2017) and the more conservative threshold of 3.3 proposed by Hair et al. (2022), with the highest observed value being 2.107, recorded for NEED_UNIQ. Consequently, the absence of critical collinearity confirms that the structural path coefficients can be interpreted as unbiased estimates of the hypothesized relationships (Hair et al., 2017). Accordingly, the evaluation of the structural model proceeds to the assessment of path coefficients in Section 5.3.2.

5.3.2 Path Coefficients and Direct Effects

To evaluate the significance, direction, and relevance of the hypothesized structural relationships in the present study, an assessment of the path coefficients is required (Hair et al., 2017). In PLS-SEM, each path coefficient is reported as a standardized beta coefficient (β) and interpreted analogously to a standardized β in OLS regression (Hair et al., 2017). Assessing these parameters requires evaluating three complementary measures: statistical significance, practical relevance, and directional consistency (Hair et al., 2017). Statistical significance refers to whether the coefficient differs from zero in the population; practical relevance addresses whether a statistically significant coefficient is large enough to carry meaningful theoretical implications; directional consistency refers to whether the sign of each significant coefficient aligns with the theoretically predicted direction (Hair et al., 2017).

Since PLS-SEM is nonparametric and distribution-free in nature, the significance of all direct path coefficients is assessed via bootstrapping with 5,000 subsamples, as established in Section 5.1 (Hair et al., 2017). Accordingly, BCa bootstrap confidence intervals are employed to construct 95% confidence intervals (Hair et al., 2017). A path coefficient is considered statistically significant when its BCa bootstrapped confidence interval does not include zero, and a 5% significance level criterion is adopted (Hair et al., 2017). Furthermore, the same bootstrapping procedure is applied to assess the specific indirect effects and the index of moderated mediation in Sections 5.3.4 and 5.3.5 respectively. The direct path coefficients are reported in Table 5.8.

Table 5.8. Direct Path Coefficients (β) of the Structural Model (Source: SmartPLS 4).

Constructs	BRAND_AUTH	MOR_DISG	PER_EFF	WTP
AI_Disclosure	-0.297	0.333	-0.835	0.011
LUX_INT	0.075			0.123
MOR_DISG	-0.170			-0.006
NEED_UNIQ		-0.118	0.038	
PER_EFF	0.233			0.162
VIS_APP	0.225			-0.049
NEED_UNIQ x AI_Disclosure		-0.110	0.023	

Note: Bootstrapped significance for the direct stimulus-response paths: AI_Disclosure $\rightarrow$ BRAND_AUTH ($t = 1.975$, $p = 0.048$); AI_Disclosure $\rightarrow$ WTP ($t = 0.063$, $p = 0.950$). All remaining inferential statistics are reported in Table 5.9.

As detailed in Table 5.8, the structural model's initial stimulus-organism ($S \rightarrow O$) pathways reveal distinct directional effects. Notably, the path from AI_Disclosure to PER_EFF yields a coefficient of $\beta = -0.835$. This represents a strong negative relationship that is directionally consistent with the hypothesized H2a in Chapter 3, and is also a large coefficient by conventional PLS-SEM standards (Hair et al., 2017). Furthermore, the path from AI_Disclosure to MOR_DISG has a coefficient of $\beta = 0.333$, demonstrating a positive direction consistent with hypothesis H3a. The moderating framework is assessed without drawing inferential conclusions, revealing that the NEED_UNIQ main effect yields a positive coefficient on PER_EFF ($\beta = 0.038$) and a negative coefficient on MOR_DISG ($\beta = -0.118$). The associated interaction term, NEED_UNIQ $\times$ AI_Disclosure, exhibits a positive coefficient on PER_EFF ($\beta = 0.023$) and a negative coefficient on MOR_DISG ($\beta = -0.110$). Both the main effects of NEED_UNIQ and the associated interaction term represent necessary model components arising from the inclusion of the moderator, rather than directly hypothesized relationships. In the organism-response pathway ($O \rightarrow R$), the parameter estimates are evaluated for the two parallel mediating constructs. The paths from PER_EFF have coefficients of $\beta = 0.162$ for PER_EFF $\rightarrow$ WTP and $\beta = 0.233$ for PER_EFF $\rightarrow$ BRAND_AUTH. Both paths demonstrate a positive direction that is directionally consistent with hypothesis H2b. Conversely, the paths originating from MOR_DISG generate negative coefficients, with a $\beta = -0.006$ for MOR_DISG $\rightarrow$ WTP and a β

= -0.170 for MOR_DISG → BRAND_AUTH. While both paths are directionally consistent with the negative relationships hypothesized in H3b, the coefficient for MOR_DISG → WTP is close to zero, suggesting that Moral Disgust exerts a negligible direct influence on WTP (Hair et al., 2017). The evaluation of the direct stimulus-response ($S \rightarrow R$) shows that the path from AI_Disclosure to BRAND_AUTH produces a coefficient of $\beta = -0.297$, which the bootstrapping procedure confirms as statistically significant ($t = 1.975$, $p = 0.048$). Since hypothesis H1b concerns the total effect of AI disclosure on Brand Authenticity rather than the isolated direct path, the formal evaluation of hypothesis H1b is conducted via the bootstrapped total effects reported in Table 5.9. Similarly, the direct AI_Disclosure → WTP path yields a coefficient of $\beta = 0.011$, which is non-significant ($t = 0.063$, $p = 0.950$). Since hypothesis H1a concerns the total effect of AI disclosure on WTP rather than the isolated direct path, this near-zero positive direct coefficient does not constitute a test of hypothesis H1a. The formal evaluation of hypothesis H1a is conducted via the bootstrapped total effects reported in Table 5.9. The structural model further includes four control variable paths, producing coefficients of $\beta = 0.123$ for LUX_INT → WTP, $\beta = 0.075$ for LUX_INT → BRAND_AUTH, $\beta = -0.049$ for VIS_APP → WTP, and $\beta = 0.225$ for VIS_APP → BRAND_AUTH.

Having reported the β estimates for all structural paths, the analysis proceeds to the assessment of all structural paths via BCa bootstrapped confidence intervals. As established in Section 5.1, the significance of all structural path coefficients is evaluated via a bootstrapping procedure, which generates empirical t-statistics and p-values for each path coefficient (Hair et al., 2017). A two-tailed test is applied, for which the critical t-value at the 5% significance level is 1.96, corresponding to a p-value threshold of 0.05 (Hair et al., 2017). Statistical significance is primarily assessed through p-values, such that structural relationships are generally considered significant when $p < 0.05$ (Hair et al., 2017). The total effects and their bootstrapped inferential results are reported in Table 5.9, while the corresponding BCa bootstrapped confidence intervals are reported in Table 5.10. For paths without mediating constructs (specifically the $O \rightarrow R$ paths and control variable paths), total effects and direct path coefficients are equivalent. For AI_Disclosure paths to WTP and BRAND_AUTH, the total effects comprise both direct and indirect contributions, as detailed in Section 5.3.4.

Table 5.9. Bootstrapping Total Effects results (Source: SmartPLS 4).

Path	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values
AI_Disclosure -> BRAND_AUTH	-0.549	-0.545	0.139	3.958	0.000
AI_Disclosure -> MOR_DISG	0.333	0.350	0.139	2.400	0.016
AI_Disclosure -> PER_EFF	-0.835	-0.827	0.126	6.640	0.000
AI_Disclosure -> WTP	-0.127	-0.139	0.163	0.778	0.437
LUX_INT -> BRAND_AUTH	0.075	0.074	0.072	1.036	0.300
LUX_INT -> WTP	0.123	0.124	0.070	1.753	0.080
MOR_DISG -> BRAND_AUTH	-0.170	-0.180	0.077	2.207	0.027
MOR_DISG -> WTP	-0.006	-0.001	0.100	0.057	0.954
NEED_UNIQ -> MOR_DISG	-0.118	-0.106	0.109	1.082	0.279
NEED_UNIQ -> PER_EFF	0.038	0.030	0.110	0.349	0.727
NEED_UNIQ x AI_Disclosure -> MOR_DISG	-0.110	-0.098	0.234	0.468	0.640
NEED_UNIQ x AI_Disclosure -> PER_EFF	0.023	0.007	0.161	0.143	0.887
PER_EFF -> BRAND_AUTH	0.233	0.233	0.083	2.807	0.005
PER_EFF -> WTP	0.162	0.162	0.068	2.400	0.016
VIS_APP -> BRAND_AUTH	0.225	0.229	0.091	2.475	0.013
VIS_APP -> WTP	-0.049	-0.035	0.101	0.482	0.630

Note: p-values reported as 0.000 indicate $p < 0.001$.

Table 5.10. *BCa Bootstrapped Confidence Intervals for the Total Effects of the Structural Model*
(Source: SmartPLS 4).

Path	Original sample (O)	Sample mean (M)	Bias	2.5%	97.5%
AI_Disclosure -> BRAND_AUTH	-0.549	-0.545	0.003	-0.806	-0.275
AI_Disclosure -> MOR_DISG	0.333	0.350	0.017	0.018	0.547
AI_Disclosure -> PER_EFF	-0.835	-0.827	0.008	-1.064	-0.571
AI_Disclosure -> WTP	-0.127	-0.139	-0.013	-0.444	0.177
LUX_INT -> BRAND_AUTH	0.075	0.074	-0.000	-0.061	0.220
LUX_INT -> WTP	0.123	0.124	0.000	-0.023	0.255
MOR_DISG -> BRAND_AUTH	-0.170	-0.180	-0.010	-0.303	-0.005
MOR_DISG -> WTP	-0.006	-0.001	0.004	-0.141	0.270
NEED_UNIQ -> MOR_DISG	-0.118	-0.106	0.012	-0.395	0.048
NEED_UNIQ -> PER_EFF	0.038	0.030	-0.008	-0.168	0.261
NEED_UNIQ x AI_Disclosure - > MOR_DISG	-0.110	-0.098	0.011	-0.484	0.414
NEED_UNIQ x AI_Disclosure - > PER_EFF	0.023	0.007	-0.016	-0.277	0.348
PER_EFF -> BRAND_AUTH	0.233	0.233	-0.000	0.060	0.385
PER_EFF -> WTP	0.162	0.162	0.000	0.014	0.282
VIS_APP -> BRAND_AUTH	0.225	0.229	0.004	0.028	0.391
VIS_APP -> WTP	-0.049	-0.035	0.014	-0.231	0.152

Note: The Original Sample (O) values correspond to the total effects reported in Table 5.9. Confidence intervals derived from bootstrapping with 5,000 subsamples (Hair et al., 2017). Paths with confidence intervals excluding zero are considered statistically significant at $p < 0.05$. p -values reported as 0.000 indicate $p < 0.001$.

As reported in Table 5.10, the BCa bootstrapped confidence intervals confirm the significance decisions across all structural paths: confidence intervals exclude zero for all paths reaching statistical significance at $p < 0.05$, while they include zero for all non-significant paths, consistent with the p -value-based conclusions reported in Table 5.9 (Hair et al., 2017).

Beginning with the $S \rightarrow O$ pathway, the bootstrapped analysis reported in Table 5.9 reveals that the direct effect of AI_Disclosure on PER_EFF yields a path coefficient of $\beta = -0.835$ ($t = 6.640$, $p < 0.001$), confirming support for hypothesis H2a. As established in Chapter 3, this finding is consistent with the effort heuristic, whereby AI disclosure signals automated production, reducing perceived effort attributions. Furthermore, the direct effect of AI_Disclosure on MOR_DISG is positive and significant ($\beta = 0.333$, $t = 2.400$, $p = 0.016$), making hypothesis H3a supported. This outcome is consistent with the moral disgust mechanism proposed in Chapter 3, whereby AI disclosure within a luxury craftsmanship context triggers a visceral affective reaction of moral violation in consumers.

The moderated relationships proposed by H4a and H4b are both non-significant. Specifically, the interaction term $NEED_UNIQ \times AI_Disclosure$ exerts no significant effect on PER_EFF ($\beta = 0.023$, $t = 0.143$, $p = 0.887$), making hypothesis H4a not supported. Moreover, while analytically inconclusive given the non-significance, the reverse direction warrants acknowledgement since the obtained coefficient ($\beta = 0.023$) is not only non-significant but also contrary in direction to the negative interaction effect proposed by hypothesis H4a in Chapter 3. Similarly, the interaction term exerts no significant effect on MOR_DISG ($\beta = -0.110$, $t = 0.468$, $p = 0.640$), making hypothesis H4b not supported. Again, the direction of the relationship is reversed, as the obtained coefficient ($\beta = -0.110$) is contrary in direction to the positive interaction effect proposed by H4b in Chapter 3. Moreover, the $NEED_UNIQ$ main effects also produce non-significant results for both PER_EFF ($\beta = 0.038$, $t = 0.349$, $p = 0.727$) and MOR_DISG ($\beta = -0.118$, $t = 1.082$, $p = 0.279$).

Continuing along the $O \rightarrow R$ pathway, the analysis evaluates the influence of the two mediating constructs on the two response constructs. For $PER_EFF \rightarrow WTP$, $PER_EFF \rightarrow BRAND_AUTH$, $MOR_DISG \rightarrow WTP$, and $MOR_DISG \rightarrow BRAND_AUTH$, the total effects reported in Table 5.9 are equal to the direct path coefficients reported in Table 5.8. These results reflect the absence of additional mediating constructs between the predictors and the response constructs (Hair et al., 2017). The paths originating from PER_EFF demonstrate significant positive effects on both WTP ($\beta = 0.162$, $t = 2.400$, $p = 0.016$) and BRAND_AUTH ($\beta = 0.233$, $t = 2.807$, $p = 0.005$), making hypothesis H2b supported. These

findings are consistent with the effort heuristic established in Chapter 3, whereby the cognitive perception of lower effort (signaled by AI disclosure) undermines the justification for the brand's premium price positioning, reducing both consumers' WTP and their perceived Brand Authenticity. Nevertheless, the evaluation of MOR_DISG reveals divergent outcomes across the two response variables. The hypothesized negative effect of MOR_DISG on WTP is not significant ($\beta = -0.006$, $t = 0.057$, $p = 0.954$). However, the corresponding negative effect on BRAND_AUTH is statistically significant ($\beta = -0.170$, $t = 2.207$, $p = 0.027$). Consequently, hypothesis H3b is partially supported: while the predicted negative effect on Brand Authenticity is confirmed, the corresponding effect on WTP is not significant. Therefore, the divergence between the significant effect on BRAND_AUTH and the non-significant effect on WTP suggests that the emotional response of Moral Disgust primarily damages consumers' perception of Brand Authenticity, as discussed in Chapter 3, rather than directly lowering their WTP.

To capture the overall economic and psychological consequences of the experimental stimulus, the total $S \rightarrow R$ effects of AI_Disclosure are also examined. The total effect of AI_Disclosure on BRAND_AUTH is negative and statistically significant ($\beta = -0.549$, $t = 3.958$, $p < 0.001$), making hypothesis H1b supported. This total effect is larger than the isolated direct path ($\beta = -0.297$) previously reported in Table 5.8, reflecting the mediation driven by the indirect path contributions via both PER_EFF and MOR_DISG (Hair et al., 2017). This finding is consistent with the theory proposed in Chapter 3, whereby AI disclosure in luxury advertising eliminates the expected human factor from brand communication, and damages perceived Brand Authenticity through both the cognitive devaluation of the premium value proposition and the affective moral violation triggered by the perception of inauthenticity. Conversely, while the total effect direction of AI_Disclosure on WTP ($\beta = -0.127$) is consistent with hypothesis H1a's prediction of a negative effect, the effect does not reach statistical significance ($t = 0.778$, $p = 0.437$). Consequently, hypothesis H1a is not supported, as the combined direct and indirect influences of AI disclosure on WTP do not generate a statistically significant total effect within the present sample. This result may partially reflect the relatively small sample

size of $N = 145$, which limits statistical power for detecting small-to-moderate effects (Hair et al., 2017). Notably, the direction of the total effect ($\beta = -0.127$) remains consistent with the predicted negative direction, suggesting the relationship may exist but remain statistically undetected at current power levels. Furthermore, the positive direct path ($\beta = 0.011$, Table 5.8) and the negative total effect ($\beta = -0.127$, Table 5.9) reveal a direction inversion, suggesting that the mediation via PER_EFF and MOR_DISG not only contributes to the total effect influence but reverses its direction relative to the isolated direct path. The implications of this result are examined in Chapter 6.

Lastly, the assessment accounts for the model's control variables. The paths from LUX_INT exhibit no significant effects on either WTP ($\beta = 0.123$, $t = 1.753$, $p = 0.080$) or BRAND_AUTH ($\beta = 0.075$, $t = 1.036$, $p = 0.300$). Although the path from LUX_INT to WTP is not statistically significant, it approaches the operative 5% significance threshold, and this warrants further examination in future research with larger samples. Furthermore, while the effect of VIS_APP on WTP is not significant ($\beta = -0.049$, $t = 0.482$, $p = 0.630$), its corresponding effect on BRAND_AUTH emerges as positive and statistically significant ($\beta = 0.225$, $t = 2.475$, $p = 0.013$). This confirms Visual Appeal as a theoretically plausible effect in luxury advertising evaluations, and its inclusion in the structural model ensures that the primary effects of AI_Disclosure on Brand Authenticity are estimated without differences in quality perception.

Following Hair et al. (2017), path coefficients below 0.10 are considered weak, those between 0.10 and 0.30 moderate, and those above 0.30 substantial. Accordingly, the significant paths in the present model range from moderate to substantial. The substantial paths include AI_Disclosure $\rightarrow$ PER_EFF ($\beta = -0.835$), AI_Disclosure $\rightarrow$ MOR_DISG ($\beta = 0.333$), and AI_Disclosure $\rightarrow$ BRAND_AUTH total effect ($\beta = -0.549$), while the moderate paths comprise PER_EFF $\rightarrow$ WTP ($\beta = 0.162$), PER_EFF $\rightarrow$ BRAND_AUTH ($\beta = 0.233$), MOR_DISG $\rightarrow$ BRAND_AUTH ($\beta = -0.170$), and VIS_APP $\rightarrow$ BRAND_AUTH ($\beta = 0.225$), confirming that all supported relationships carry meaningful practical relevance beyond statistical significance (Hair et al., 2017).

Taken together, of the eight hypotheses tested, four are supported (H1b, H2a, H2b, H3a), one is partially supported (H3b), and three are not supported (H1a, H4a, H4b). The hypotheses concerning moderated mediation (H4c and H4d) are evaluated in Section 5.3.5. Having evaluated the significance and directional hypotheses for all structural relationships across the specified pathways, a summary of all hypotheses is provided in Section 5.5. Next, the assessment of the model's explanatory power is addressed.

5.3.3 Model's Explanatory Power and Effect Size (R^2 , f^2)

Having established the significance, direction, and practical relevance of the structural path coefficients in Section 5.3.2, the analysis proceeds to evaluate the overall explanatory power of the structural model. Following the procedure established by Hair et al. (2017), this assessment is conducted through two complementary criteria: the Coefficient of Determination (R^2) and the Effect Size (f^2).

The coefficient of determination (R^2) is defined as the proportion of variance in each endogenous construct explained by its predictor constructs in the structural model (Hair et al., 2017). R^2 is calculated as the squared correlation between the actual and the predicted values of the endogenous construct, ranging from 0 to 1, where higher values indicate greater explanatory power (Hair et al., 2017). In PLS-SEM, R^2 is computed separately for each endogenous construct given the model's multiple regression structure, as established in Section 5.3.1. The interpretation of R^2 values generally follow the benchmarks established in marketing research literature. Accordingly, R^2 values of 0.75, 0.50, and 0.25 are typically interpreted as substantial, moderate, and weak respectively, providing a standard that is considered appropriate for evaluating the model's explanatory power (Hair et al., 2017). Nevertheless, acceptable R^2 values tend to vary by research context and model complexity, such that more parsimonious models or studies examining complex behavioral phenomena may yield lower R^2 . Consequently, the adjusted R^2 (R^2_{adj}) is introduced as a complementary measure that corrects for model complexity by penalizing the inclusion of additional predictor constructs that do not contribute meaningfully to explanatory power (Hair et al., 2017). R^2_{adj} provides a more

conservative and appropriate benchmark in models with multiple predictors (Hair et al., 2017). The R^2 and R^2_{adj} values for all endogenous constructs in the present structural model are reported in Table 5.11.

Table 5.11. R^2 and R^2_{adj} for the Structural Model (Source: SmartPLS 4).

Construct	R^2	R^2_{adj}	Interpretation
BRAND_AUTH	0.234	0.207	Below Weak
MOR_DISG	0.056	0.036	Very Low
PER_EFF	0.172	0.154	Low
WTP	0.038	0.003	Very Low

As referenced in Table 5.11, for BRAND_AUTH, the convergence of predictors (AI_Disclosure, PER_EFF, MOR_DISG, VIS_APP, and LUX_INT, of which the former four exert significant effects as established in Section 5.3.2) collectively accounts for 23.4% of its variance ($R^2 = 0.234$, $R^2_{adj} = 0.207$), falling just below the weak threshold of 0.25 established by Hair et al. (2017). Moreover, for PER_EFF ($R^2 = 0.172$, $R^2_{adj} = 0.154$), AI_Disclosure alone accounts for 17.2% of the variance, reflecting the strong influence of the experimental stimulus on the cognitive pathway, given that NEED_UNIQ and the interaction term were both non-significant, as established in Section 5.3.2. Conversely, AI_Disclosure alone accounts for only 5.6% of the variance in MOR_DISG ($R^2 = 0.056$, $R^2_{adj} = 0.036$), suggesting substantial unexplained variance in the affective pathway. Similarly, only PER_EFF contributes significantly to WTP ($R^2 = 0.038$, $R^2_{adj} = 0.003$), accounting for the lowest explained variance in the model (3.8%). Notably, the R^2 values for WTP and MOR_DISG fall well below the weak threshold of 0.25 established by Hair et al. (2017), representing a limitation of the present model in fully accounting for the variance in these two constructs. Nevertheless, this outcome is expected given the experimental nature of the present study, as a single binary manipulation (human photography versus AI-generated image condition) can only account for a limited portion of the total variance in consumer responses (Hair et al., 2017). Furthermore, the gap between R^2 and R^2_{adj} for

WTP is larger than the remaining constructs, indicating that its predictor set adds limited explanatory value, and suggesting that WTP is driven by factors beyond those captured in the present model. Accordingly, the individual contribution of each predictor to R^2 is examined through f^2 effect sizes.

The f^2 effect size is a complementary measure that, unlike R^2 which captures the overall explanatory power of all predictors combined, isolates the individual contribution of each specific exogenous construct to the R^2 of its corresponding endogenous construct (Hair et al., 2017). This is generally calculated by comparing the R^2 value of the endogenous construct with and without the predictor of interest, thereby quantifying the change in explanatory power attributable to the removal of that specific predictor (Hair et al., 2017). The f^2 assessment is conducted independently of whether individual path coefficients reach statistical significance, as a predictor may contribute meaningfully to explanatory power despite a non-significant path, and conversely, a statistically significant path may yield a negligible f^2 (Hair et al., 2017). The interpretation of these parameters follows the established benchmarks proposed by Cohen (1988) and adopted by Hair et al. (2017), whereby f^2 values of 0.02, 0.15, and 0.35 indicate small, medium, and large effects respectively, while values falling below 0.02 suggest a largely negligible contribution. Accordingly, the f^2 values for all structural paths in the present model are reported in Table 5.12.

Table 5.12. f^2 for the Structural Model (Source: SmartPLS 4).

Construct	BRAND_AUTH	MOR_DISG	PER_EFF	WTP
AI_Disclosure	0.022	0.029	0.207	0.000
LUX_INT	0.007			0.015
MOR_DISG	0.032			0.000
NEED_UNIQ		0.007	0.001	
PER_EFF	0.057			0.022
VIS_APP	0.057			0.002
NEED_UNIQ x AI_Disclosure		0.003	0.000	

As referenced in Table 5.12, for the endogenous construct PER_EFF, AI_Disclosure yields an f^2 value of 0.207, which is classified as a medium-to-large effect (Cohen, 1988; Hair et al., 2017). This is the single medium-to-large effect within the entire model, a finding consistent with the large path coefficient ($\beta = -0.835$) reported in Section 5.3.2, confirming the disproportionate influence of the experimental stimulus on the cognitive pathway. Conversely, NEED_UNIQ ($f^2 = 0.001$) and the interaction term NEED_UNIQ $\times$ AI_Disclosure ($f^2 = 0.003$) yield negligible contributions, supporting the non-significant path coefficients reported in Section 5.3.2. Furthermore, regarding MOR_DISG, AI_Disclosure exerts a small effect ($f^2 = 0.029$), consistent with the significant but moderate path coefficient ($\beta = 0.333$), while NEED_UNIQ ($f^2 = 0.007$) and the interaction term ($f^2 = 0.000$) yield negligible contributions.

For the response constructs, the assessment of BRAND_AUTH indicates that no single predictor dominates. PER_EFF ($f^2 = 0.057$) and VIS_APP ($f^2 = 0.057$) share the largest effects, both classified as small, followed by MOR_DISG ($f^2 = 0.032$, small) and AI_Disclosure ($f^2 = 0.022$, small), whereas LUX_INT yields a negligible contribution ($f^2 = 0.007$). Notably, VIS_APP $\rightarrow$ BRAND_AUTH is statistically significant ($p = 0.013$) yet yields only a small f^2 of 0.057, confirming that statistical significance and practical relevance do not always converge (Cohen, 1988; Hair et al., 2017). These small effects across multiple predictors collectively account for the R^2 of 0.234 reported in Table 5.11, which falls just below the weak threshold despite the convergence of four significant predictors. The analysis of WTP reveals that only PER_EFF yields a small effect ($f^2 = 0.022$). Nevertheless, LUX_INT ($f^2 = 0.015$) approaches but does not reach the small effect threshold of 0.02, a pattern that mirrors the p-value of 0.080 reported in Section 5.3.2. Both significance and effect size criteria therefore converge in indicating a marginal relationship that warrants further investigation in future research with larger samples (Hair et al., 2017). Conversely, AI_Disclosure ($f^2 = 0.000$) and MOR_DISG ($f^2 = 0.000$) both yield negligible values, indicating that neither the experimental stimulus nor the affective mediator contributes to explaining WTP variance beyond what other predictors account for. VIS_APP similarly yields a negligible contribution ($f^2 = 0.002$), further confirming that WTP remains largely unaccounted for by the present set of predictors.

Consequently, the f^2 results collectively confirm the pattern established in Section 5.3.2. The cognitive pathway AI_Disclosure → PER_EFF is the primary driver of mediator-level variance, while the affective pathway via MOR_DISG exerts selective influence primarily on BRAND_AUTH (consistent with the partial support of H3b established in Section 5.3.2). At the response level, no single predictor dominates, as the explanatory power of BRAND_AUTH reflects the cumulative contribution of multiple small effects rather than a single dominant pathway. Furthermore, the moderation contributes negligibly to the model's explanatory power. Accordingly, the analysis proceeds to the parallel mediation assessment in Section 5.3.4.

5.3.4 Parallel Mediation Analysis: Perceived Effort and Moral Disgust

According to Hair et al. (2017), mediation represents the mechanism through which the relationship between an exogenous and an endogenous construct is transmitted by one or more intervening variables, such that a change in the exogenous construct produces a variation in the mediating construct, which in turn influences the endogenous construct. Generally, incorporating mediating constructs enables a more precise understanding of the underlying psychological mechanisms through which the stimulus exerts its effect on consumer responses (Hair et al., 2017). The incorporation of mediation into the analytical framework requires a solid theoretical foundation, as established in Chapters 2 and 3, to ensure validity and interpretability of empirical results. In the present model, AI_Disclosure operates as the exogenous stimulus, transmitting its effects on WTP and BRAND_AUTH through two simultaneous parallel mediating pathways: Perceived Effort (PER_EFF) representing the cognitive pathway and Moral Disgust (MOR_DISG) representing the affective pathway, as illustrated in Figure 5.2 and theoretically established in Chapter 3.

The assessment of mediation effects requires that both the measurement model and the structural model satisfy the established reliability, validity, and collinearity criteria (conditions confirmed in Sections 5.2 and 5.3.1 respectively) (Hair et al., 2017). Furthermore, the present analysis applies the mediation classification framework established by Zhao et al. (2010) and adopted by Hair et al. (2017), which organizes

mediation outcomes into five mutually exclusive categories divided into two groups. The first group is characterized by the absence of mediation and includes direct-only non-mediation, in which the direct effect is significant but the indirect effect is not, and no-effect non-mediation, in which neither the direct nor the indirect effect reaches statistical significance (Hair et al., 2017; Zhao et al., 2010). The second group is characterized by the presence of a mediating effect and includes: complementary mediation, in which both the indirect and direct effects are significant and operate in the same direction; competitive mediation, in which both effects are significant but operate in opposite directions; and indirect-only mediation, in which the indirect effect is significant but the direct effect is not (Hair et al., 2017; Zhao et al., 2010). A related classification system, proposed in the PLS-SEM Handbook of Esposito Vinzi et al. (2010), similarly distinguishes between full mediation (corresponding to indirect-only mediation) and partial mediation (corresponding to complementary or competitive mediation), and no mediation (corresponding to direct-only and no-effect non-mediation). While the two frameworks are equivalent in their logic, the present study adopts the Zhao et al. (2010) terminology. Accordingly, the direct effect of AI_Disclosure on each response construct serves as the reference against which each indirect effect is evaluated for mediation classification (Hair et al., 2017; Zhao et al., 2010). When collinearity among predictors is high, the signs of direct and indirect effects may be distorted, potentially generating artificially opposite signs (Hair et al., 2017). Nevertheless, given that all the VIF values fell below the operative threshold of 5.0 (confirmed in Section 5.3.1), this concern does not apply to the present study.

Consequently, the assessment of the parallel mediation model follows the procedure established by Hair et al. (2017), which requires evaluating the significance of specific indirect effects as the primary step. If the indirect path does not reach statistical significance, mediation cannot be supported regardless of the direct effect outcome (Hair et al., 2017). Specific indirect effects are computed as the product of the path coefficients constituting each mediation chain specified in the present model. Furthermore, their significance is assessed via BCa bootstrapped confidence intervals with 5,000

subsamples, as noted in Section 5.1. Accordingly, the specific indirect effects and their bootstrapped inferential results are reported in Table 5.13. Table 5.14 reports the confidence intervals bias corrected.

Table 5.13. *Specific Indirect Effects of the Structural Model (Source: SmartPLS 4).*

Mediation Chain	Sign	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values
AI_Disclosure -> PER_EFF -> WTP	-	-0.136	-0.135	0.063	2.157	0.031
AI_Disclosure -> MOR_DISG -> WTP	-	-0.002	0.002	0.038	0.051	0.960
AI_Disclosure -> PER_EFF -> BRAND_AUTH	-	-0.195	-0.192	0.074	2.634	0.008
AI_Disclosure -> MOR_DISG -> BRAND_AUTH	-	-0.057	-0.063	0.037	1.522	0.128

Table 5.14. *BCa Bootstrapped Confidence Intervals for the Specific Indirect Effects of the Structural Model (Source: SmartPLS 4).*

Mediation Chain	Original sample (O)	Sample mean (M)	Bias	2.5%	97.5%
AI_Disclosure -> PER_EFF -> WTP	-0.136	-0.135	0.001	-0.270	-0.020
AI_Disclosure -> MOR_DISG -> WTP	-0.002	0.002	0.004	-0.053	0.106
AI_Disclosure -> PER_EFF -> BRAND_AUTH	-0.195	-0.192	0.003	-0.354	-0.060
AI_Disclosure -> MOR_DISG -> BRAND_AUTH	-0.057	-0.063	-0.006	-0.134	0.003

As reported in Tables 5.13 and 5.14, the specific indirect effect of AI_Disclosure on WTP through PER_EFF is statistically significant ($\beta = -0.136$, $t = 2.157$, $p = 0.031$), as is the corresponding specific indirect effect on BRAND_AUTH ($\beta = -0.195$, $t = 2.634$, $p = 0.008$) and, for both effects, the BCa bootstrapped confidence intervals exclude zero. Accordingly, both specific indirect effects are significant and directionally consistent with the predicted negative sign, making hypothesis H2 supported. The specific indirect effect of AI_Disclosure on WTP through MOR_DISG yields a coefficient of $\beta = -0.002$ ($t = 0.051$, $p = 0.960$), while the specific indirect effect on BRAND_AUTH produces a coefficient of $\beta = -0.057$ ($t = 1.522$, $p = 0.128$). Consequently, neither effect reaches statistical significance at $p < 0.05$, making hypothesis H3 not supported. Furthermore, the BCa bootstrapped confidence intervals (Table 5.14) for both indirect effects include zero, confirming the non-significance of both indirect effects and establishing that the affective pathway does not transmit the effect of AI_Disclosure on either response construct (Hair et al., 2017). Nevertheless, while both indirect effects are directionally consistent with the predicted negative sign, suggesting the mechanism operates in the hypothesized direction, the near-zero value for MOR_DISG $\rightarrow$ WTP ($\beta = -0.002$) indicates a negligible indirect effect on WTP, while MOR_DISG $\rightarrow$ BRAND_AUTH ($\beta = -0.057$) approaches but does not reach the significance threshold. The non-significance of the MOR_DISG $\rightarrow$ BRAND_AUTH indirect effect may partially reflect the relatively small sample size of $N = 145$, limiting statistical power for detecting small indirect effects (Hair et al., 2017).

Following the mediation classification framework of Zhao et al. (2010) adopted by Hair et al. (2017), the cognitive pathway yields two distinct outcomes. The AI_Disclosure $\rightarrow$ PER_EFF $\rightarrow$ WTP chain is classified as indirect-only mediation, as the specific indirect effect is significant while the direct effect of AI_Disclosure on WTP ($\beta = 0.011$, $p = 0.950$) is not. This outcome underscores that the negative influence of AI disclosure on consumers' WTP does not operate independently of the cognitive pathway, but it depends entirely on the cognitive pathway. AI disclosure reduces perceived craft, and it is this reduction in perceived craft that in turn lowers consumers' WTP. Without the mediation of PER_EFF, no significant direct economic consequence of the disclosure is detected.

Under the Zhao et al. (2010) framework, this represents the theoretically ideal mediation scenario, as the cognitive mediator fully accounts for the stimulus-to-response transmission.

The AI_Disclosure → PER_EFF → BRAND_AUTH chain is classified as complementary mediation, as both the specific indirect effect (Table 5.13) and the direct effect of AI_Disclosure on BRAND_AUTH ($\beta = -0.297$, $p = 0.048$) are statistically significant and operate in the same negative direction. In the present study, this means that PER_EFF partially, but not fully, explains why AI disclosure damages Brand Authenticity evaluations, and an independent residual negative effect persists. According to Zhao et al. (2010), complementary mediation of this kind indicates that an additional mediating mechanism may be present.

The AI_Disclosure → MOR_DISG → WTP chain is classified as no-effect non-mediation, as neither the specific indirect effect nor the direct effect of AI_Disclosure on WTP ($\beta = 0.011$, $p = 0.950$) reaches statistical significance, meaning that Moral Disgust does not transmit the effect of the AI disclosure on consumers' WTP, and no independent direct relationship between the stimulus and the WTP is empirically supported either. Accordingly, Hair et al. (2017) note that this may indicate a flawed theoretical framework, suggesting that the affective route connecting AI disclosure to WTP via Moral Disgust may require further theoretical discussion.

Lastly, the AI_Disclosure → MOR_DISG → BRAND_AUTH chain is classified as direct-only non-mediation, as the direct effect of AI_Disclosure on BRAND_AUTH ($\beta = -0.297$, $p = 0.048$) is statistically significant while the specific indirect pathway through MOR_DISG does not reach significance. In the present study, this means that while MOR_DISG exerts a confirmed direct influence on BRAND_AUTH, it does not transmit the effect of AI_Disclosure on brand authenticity evaluations as a mediating construct (Hair et al., 2017; Zhao et al., 2010). Therefore, the AI label appears to damage brand authenticity through a mechanism that bypasses the Moral Disgust, suggesting that an alternative affective pathway may be present. The theoretical implications of the mediation classification findings are addressed in Chapter 6.

Following the specific indirect effect assessment, the analysis proceeds to examine the total indirect effects of AI_Disclosure on WTP and BRAND_AUTH. In the context of a multiple-mediation model, Hair et al. (2017) establish that the assessment of total indirect effects is a necessary complement to the specific indirect effect analysis. Total indirect effects capture the combined contribution of both mediating constructs to the overall indirect influence of AI_Disclosure on each response construct. Accordingly, the total indirect effects and their bootstrapped inferential results are reported in Table 5.15.

Table 5.15. Total Indirect Effect for the Structural Model (Source: SmartPLS 4).

Mediation Path	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values
AI_Disclosure -> BRAND_AUTH	-0.252	-0.255	0.085	2.977	0.003
AI_Disclosure -> WTP	-0.137	-0.133	0.067	2.059	0.040

As referenced in Table 5.15, the total indirect effect of AI_Disclosure on BRAND_AUTH is negative and statistically significant ($\beta = -0.252$, $t = 2.977$, $p = 0.003$), representing the combined indirect contribution of both PER_EFF and MOR_DISG. Moreover, this total indirect effect complements the total effect reported in Table 5.9, accounting for approximately 45.9% of the total effect of AI_Disclosure on BRAND_AUTH ($\beta = -0.549$, Table 5.9). Furthermore, the total indirect effect of AI_Disclosure on WTP is also statistically significant ($\beta = -0.137$, $t = 2.059$, $p = 0.040$). Moreover, the BCa bootstrapped confidence intervals exclude zero for both total indirect effects, further confirming their significance (Hair et al., 2017). Nevertheless, while the indirect pathway transmits a significant negative effect on WTP, the positive direct path from AI_Disclosure to WTP ($\beta = 0.011$, Table 5.8) partially offsets this negative indirect contribution, resulting in a non-significant total effect ($\beta = -0.127$, $p = 0.437$, Table 5.9), as established in Section 5.3.2. The broader theoretical implications of these findings are addressed in Chapter 6.

Following the assessment of both specific and total indirect effects, the cognitive pathway via PER_EFF emerges as the primary mediation mechanism of the present model. Nevertheless, while the affective pathway via MOR_DISG does not produce statistically significant specific indirect effects, it contributes to the significant total indirect effect on BRAND_AUTH. Accordingly, the analysis proceeds to the moderated mediation assessment in Section 5.3.5.

5.3.5 Moderated Mediation Analysis: Need for Uniqueness

In the present study, Need for Uniqueness (NEED_UNIQ) functions as a boundary condition moderating the influence of AI_Disclosure on the two organism constructs. According to Hair et al. (2017), moderation can be defined as the condition in which a third variable, the moderator, influences the strength or direction of the relationship between an exogenous and an endogenous construct, such that the effect varies across different levels of the moderating variable. Furthermore, an extension of moderation is moderated mediation, defined as the analytical mechanism through which a moderating variable influences the strength or direction of a mediated relationship, such that the mediating pathway itself varies as a function of the moderator (Hair et al., 2017). In PLS-SEM, researchers tend to distinguish between two types of moderation. In independent variable moderation, the moderator influences the relationship between the exogenous construct and the mediator; in dependent variable moderation, the moderator influences the relationship between the mediator and the endogenous construct (Hair et al., 2017). Accordingly, the present study applies independent variable moderation, as NEED_UNIQ moderates the paths from AI_Disclosure to PER_EFF (H4a) and from AI_Disclosure to MOR_DISG (H4b). As established in Chapter 3, consumers with a higher Need for Uniqueness are theorized to be more sensitive to the standardization signal carried by AI disclosure, given that AI-generated content threatens the exclusivity and individuality associated with luxury consumption (Loureiro et al., 2024). Consequently, the index of moderated mediation, tested through H4c and H4d, assesses whether the conditional indirect effect of AI_Disclosure on WTP and BRAND_AUTH, transmitted through the mediating pathways, varies as a function of NEED_UNIQ, thereby testing whether the entire mediation chain is itself conditioned by the moderator rather than only the first-

stage paths (Hair et al., 2017; Hayes, 2015). Furthermore, the main effects of NEED_UNIQ on both PER_EFF and MOR_DISG were reported as non-significant structural components in Section 5.3.2. Accordingly, the present assessment focuses exclusively on the interaction term as a test of the moderation hypotheses.

The moderated mediation analysis is conducted using the two-stage approach established in Section 5.1, as recommended by Hair et al. (2017) for models with reflective constructs. This approach first estimates the main effect model without the interaction term to obtain latent variable scores, which are then multiplied in the second stage to create a single-item interaction term (Hair et al., 2017). Furthermore, the resulting interaction term, NEED_UNIQ $\times$ AI_Disclosure, was assessed for collinearity in Section 5.3.1, where VIF values confirmed acceptable levels. The bootstrapped path coefficients of the interaction term NEED_UNIQ $\times$ AI_Disclosure on PER_EFF and MOR_DISG are reported in Table 5.16.

Table 5.16. Path coefficients for the Moderated Mediation (Source: SmartPLS 4).

Construct		Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values
NEED_UNIQ	x					
AI_Disclosure	->	-0.110	-0.098	0.234	0.468	0.640
MOR_DISG						
NEED_UNIQ	x					
AI_Disclosure	->	0.023	0.007	0.161	0.143	0.887
PER_EFF						

As referenced in Table 5.16, both interaction term path coefficients are non-significant, making H4a and H4b not supported. Specifically, NEED_UNIQ $\times$ AI_Disclosure $\rightarrow$ PER_EFF yields $\beta = 0.023$ ($t = 0.143$, $p = 0.887$) and NEED_UNIQ $\times$ AI_Disclosure $\rightarrow$ MOR_DISG yields $\beta = -0.110$ ($t = 0.468$, $p = 0.640$), with BCa bootstrapped confidence intervals including zero for both paths (Hair et al., 2017). As established in Section 5.3.2, the obtained directions are contrary to the hypothesized directions, a reversal that warrants

acknowledgement despite being non-significant. Nevertheless, Hair et al. (2017) and Hayes (2015) recommend proceeding to the index of moderated mediation regardless of first-stage results, as it constitutes the definitive test of whether the conditional indirect effect of AI_Disclosure on WTP and BRAND_AUTH varies significantly as a function of NEED_UNIQ across its range, thereby formally testing H4c and H4d. The index of moderated mediation and its bootstrapped inferential results are reported in Table 5.17.

Table 5.17. Index of Moderated Mediation (Source: SmartPLS 4).

Moderated Mediation Chain	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values
NEED_UNIQ x AI_Disclosure -> PER_EFF -> BRAND_AUTH	0.005	0.000	0.040	0.134	0.894
NEED_UNIQ x AI_Disclosure -> MOR_DISG -> BRAND_AUTH	0.019	0.020	0.048	0.391	0.696
NEED_UNIQ x AI_Disclosure -> PER_EFF -> WTP	0.004	0.002	0.028	0.134	0.894
NEED_UNIQ x AI_Disclosure -> MOR_DISG -> WTP	0.001	0.005	0.028	0.022	0.982

As referenced in Table 5.17, the index of moderated mediation is examined across all four conditional indirect effect chains to assess whether the conditional indirect effects of AI_Disclosure on WTP and BRAND_AUTH vary as a function of NEED_UNIQ. The analysis reveals the following conditional indirect effects: NEED_UNIQ $\times$ AI_Disclosure $\rightarrow$ PER_EFF $\rightarrow$ BRAND_AUTH ($\beta = 0.005$, $t = 0.134$, $p = 0.894$), NEED_UNIQ $\times$ AI_Disclosure $\rightarrow$ PER_EFF $\rightarrow$ WTP ($\beta = 0.004$, $t = 0.134$, $p = 0.894$), NEED_UNIQ $\times$ AI_Disclosure $\rightarrow$ MOR_DISG $\rightarrow$ BRAND_AUTH ($\beta = 0.019$, $t = 0.391$, $p = 0.696$), and NEED_UNIQ $\times$ AI_Disclosure $\rightarrow$

MOR_DISG $\rightarrow$ WTP ($\beta = 0.001$, $t = 0.022$, $p = 0.982$). Consequently, none of the four chains reaches statistical significance at $p < 0.05$, making hypotheses H4c and H4d not supported. Furthermore, the BCa bootstrapped confidence intervals include zero for all four chains, further confirming non-significance (Hair et al., 2017). Notably, all four obtained coefficients are near-zero in magnitude and positive in direction (contrary to the hypothesized negative conditional indirect effects), suggesting that NEED_UNIQ not only fails to amplify the negative effects of AI_Disclosure through either the cognitive or affective pathway but does not condition these effects in any meaningful direction.

Following Hair et al. (2017), who recommend assessing the f^2 effect size of the interaction term to evaluate the practical relevance of the moderating effect, f^2 is not computed in the present study given that the complete non-significance of all interaction term paths indicates a negligible moderating effect by definition. All four conditional indirect effect chains represent no moderated mediation, as neither the first-stage interaction paths (H4a, H4b) nor the index of moderated mediation (H4c, H4d) reach statistical significance (Hair et al., 2017; Hayes, 2015). The theoretical implications of these null findings are addressed in Chapter 6. Consequently, the analysis proceeds to the assessment of the control variables in Section 5.3.6.

5.3.6 Assessment of Control Variables: Luxury Interest and Visual Appeal

In PLS-SEM structural models, control variables are generally included to account for the potential influence of external factors on the endogenous constructs, thereby ensuring that the variance explained by the primary structural relationships is not attributable to variables outside the theoretical framework (Hair et al., 2017). The inclusion of control variables does not constitute hypothesis testing, but rather a methodological safeguard against alternative explanations of the observed effects. Accordingly, significant control variable paths demonstrate that the controlled variables exert independent influence on the endogenous constructs, while non-significant paths indicate that the controlled variables do not distort the primary structural relationships (Hair et al., 2017).

Consequently, Luxury Interest (LUX_INT) and Visual Appeal (VIS_APP) are included as the two control variables in the present structural model, as discussed in Chapter 3 and operationalized in Section 4.3. Their inclusion reflects two primary sources of influence identified in the experimental context. Pre-existing luxury interest may influence WTP and BRAND_AUTH regardless of the experimental manipulation, while perceived aesthetic quality of the advertising stimulus may vary across the human and AI conditions and thereby interfere with the primary effects of AI_Disclosure on the response constructs. Specifically, LUX_INT captures the degree to which the respondent is interested in and engaged with luxury consumption, and has been operationalized as a single-item construct as established in Section 4.3, and accordingly excluded from the reflective measurement model evaluation in Section 5.2. Furthermore, VIS_APP captures the respondent's aesthetic evaluation of the advertising stimulus, and has been operationalized as a three-item reflective construct measuring perceived visual attractiveness. Both control variables predict WTP and BRAND_AUTH directly in the structural model, with their effects assessed via BCa bootstrapped confidence intervals as established in Section 5.1. The bootstrapped path coefficients and significance decisions for the two control variables are reported in Table 5.18.

Table 5.18. Path coefficients and significance of the control variables (Source: SmartPLS 4).

Path	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/ STDEV)	P values
LUX_INT -> BRAND_AUTH	0.075	0.074	0.072	1.036	0.300
LUX_INT -> WTP	0.123	0.124	0.070	1.753	0.080
VIS_APP -> BRAND_AUTH	0.225	0.229	0.091	2.475	0.013
VIS_APP -> WTP	-0.049	-0.035	0.101	0.482	0.630

As referenced in Table 5.18, the path coefficient for LUX_INT → BRAND_AUTH is not statistically significant ($\beta = 0.075$, $t = 1.036$, $p = 0.300$). This non-significance indicates that pre-existing luxury interest does not independently drive brand authenticity evaluations in the present study, suggesting that the observed effects of AI_Disclosure on BRAND_AUTH are not attributable to respondents' general Luxury Interest but rather to the specific experimental manipulation (Hair et al., 2017). Furthermore, the path coefficient for LUX_INT → WTP yields a borderline result ($\beta = 0.123$, $t = 1.753$, $p = 0.080$) that approaches, but does not reach, the operative 5% significance threshold. Consequently, this outcome appears consistent with the power limitation noted in Section 5.3.2, suggesting a marginal relationship that warrants further examination in future research with larger samples (Hair et al., 2017).

The path coefficient for VIS_APP → BRAND_AUTH is statistically significant ($\beta = 0.225$, $t = 2.475$, $p = 0.013$), confirming that visual quality perceptions independently influence brand authenticity evaluations. This result is consistent with established findings discussed in Chapter 3, whereby visual quality serves as a peripheral cue in luxury advertising evaluations (Hair et al., 2017). Moreover, this finding validates the inclusion of VIS_APP in the structural model, demonstrating that the primary effects of AI_Disclosure on BRAND_AUTH established in Section 5.3.2 reflect the experimental manipulation rather than differences in perceived image quality across the two conditions. Conversely, the path coefficient for VIS_APP → WTP is not significant ($\beta = -0.049$, $t = 0.482$, $p = 0.630$), confirming that aesthetic quality perceptions do not significantly influence WTP in the present experimental context (Hair et al., 2017).

Taken together, the control variable assessment validates the inclusion of both LUX_INT and VIS_APP in the structural model, with VIS_APP → BRAND_AUTH emerging as the only significant control path. Accordingly, a consolidated analysis of the WTP distribution is presented in Section 5.4, while a complete summary of the tested hypotheses is presented in Section 5.5.

5.4 Analysis of the Willingness to Pay distribution

In the present study, Willingness to Pay (WTP) was operationalized as a single-item open-ended response variable, following the measurement approach proposed by Fuchs et al. (2010) and described in Chapter 4. The open-ended format avoids the anchoring effects associated with scale-based formats and is particularly relevant in luxury consumption contexts where individual price thresholds tend to vary across consumers (Kapferer & Laurent, 2016; Schmidt et al., 2024). Nevertheless, an open-ended WTP question remains susceptible to hypothetical bias in the absence of a real monetary commitment (Schmidt & Bijmolt, 2020). Accordingly, the analytic sample of $N = 145$ reflects a preliminary removal of extreme and implausible responses conducted prior to the data analysis, as described in Chapter 4.

Prior to the PLS-SEM estimation, the WTP distribution was screened to identify statistical outliers using the Interquartile Range (IQR) method. The procedure yielded the following values: $Q1 = €180$, $Q3 = €500$, $IQR = €320$, and an upper fence of $€980$ (calculated as $Q3 + 1.5 \times IQR$). Nine responses exceeded this threshold, with values ranging from $€1,000$ to $€3,500$. These responses were retained within the full analytic sample following manual verification that each represented a plausible luxury price valuation, consistent with the study context. Figure 5.3 presents the frequency distribution of WTP responses across the full analytic sample of $N = 145$.

As illustrated in Figure 5.3, the frequency distribution of WTP responses exhibits a right-skewed shape, with a mean of $€415.20$ and a median of $€300$. This divergence is consistent with the concentration of responses in the lower-to-mid price range and the presence of high-value observations in the upper tail. The standard deviation of $€453.64$ further reflects the substantial dispersion in individual price perceptions across the sample. Notably, approximately 65% of respondents reported WTP values between $€101$ and $€500$, whereas 9 respondents (6.2% of the sample) reported values exceeding $€1,000$. As established in Chapter 4, WTP responses were collected with reference to MERIDIAN, a fictitious luxury brand developed for the experimental stimuli to ensure that observed WTP variation across the two conditions reflected the AI disclosure

manipulation rather than pre-existing brand-specific associations (Ali et al., 2025; Kirk & Givi, 2025). Nevertheless, the absence of authentic brand equity and real market pricing anchors may have led to response variability and reduced validity of the WTP estimates (Diallo et al., 2021). A comparative summary of descriptive statistics across the human photography and AI-generated conditions is reported in Table 5.19.

Figure 5.3. WTP distribution for N = 145 (Source: the author).

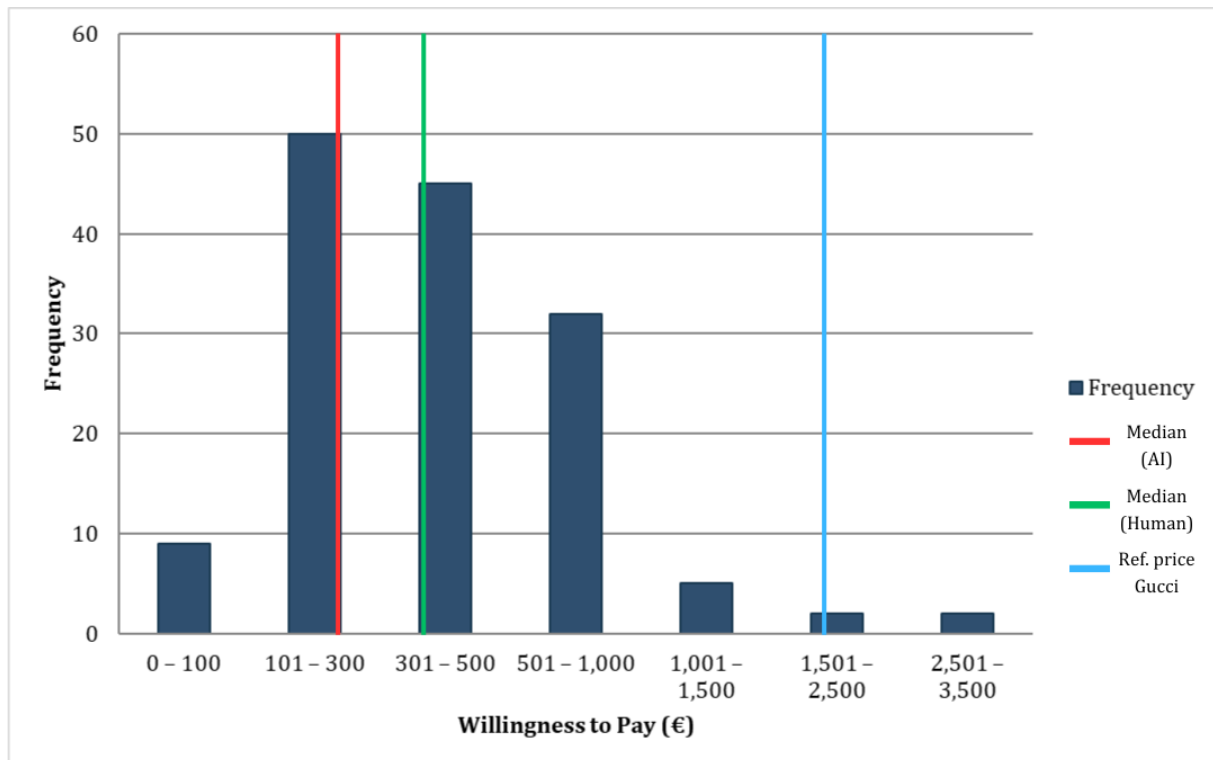


Table 5.19. WTP comparative summary across the Human and the AI-Generated conditions (Source: the author).

Condition	n	Mean	Median	STDEV
Human	65	435.60	309	437.92
AI-Generated	80	398.63	300	468.11

As detailed in Table 5.19, the Human photography condition ($n = 65$) presents a mean WTP of €435.60 and a median of €309. The AI-generated condition ($n = 80$) yields a mean of €398.63 and a median of €300, producing an absolute mean difference of approximately €37 and a median difference of €9 between the two conditions. Both measures are directionally consistent with the hypothesis that AI disclosure reduces consumers' WTP, with the Human condition yielding higher values for both the mean and the median. Furthermore, substantial dispersion is observed in both groups: the standard deviation of €437.92 in the Human condition and €468.11 in the AI-generated condition indicate that the individual price perceptions vary widely within each group, with the AI condition exhibiting marginally greater variability. This suggests that the disclosure of AI-generated content may introduce additional heterogeneity in consumers' price evaluations. Nevertheless, the differences between the groups remain modest, and no inferential comparison is conducted at this stage, as the statistical assessment of the WTP pathway has been addressed through the PLS-SEM structural model reported in Section 5.3.

The observed WTP values across both conditions must also be contextualized against current retail prices for comparable luxury weekender bags from leading luxury companies. Reference retail prices for analogous products include the Gucci Medium Duffle weekend bag at €1,650 (Gucci, 2025), the Louis Vuitton Keepall Bandoulière 45 at €2,150 (Louis Vuitton, 2025), and the Prada Re-Nylon and Saffiano leather duffle bag at €3,100 (Prada, 2025). The mean WTP values observed across both experimental conditions, ranging from approximately €399 in the AI-generated condition to €436 in the Human photography condition, fall substantially below these market reference prices. Consequently, this divergence between price valuations may be attributed to the fictitious brand context of the experimental stimuli: the absence of established brand equity, heritage, and authentic market pricing anchors associated with MERIDIAN may have suppressed the WTP relative to valuations that recognised luxury brand names carry (Diallo et al., 2021).

Taken together, the distributional characteristics reported in the present section provide a contextual foundation for interpreting the structural model findings summarized in Section 5.5. The broader theoretical and managerial implications of these distributional patterns are addressed in the general discussion in Chapter 6.

5.5 Summary of Hypothesis Testing Results

The completion of the structural model assessment across Sections 5.3.1 through 5.3.6 provides the empirical basis to evaluate the theoretical framework proposed in Chapter 3. This evaluation has addressed collinearity assessment, path coefficient significance and relevance, explanatory power, parallel mediation, moderated mediation, and the influence of control variables. The structural model assessment has therefore produced the empirical evidence necessary to evaluate each hypothesis against its theoretically derived prediction (Hair et al., 2017). Accordingly, the present section consolidates these empirical results into a unified hypothesis summary, providing a comprehensive overview of all hypothesized relationships tested in the present study.

Table 5.20 presents the consolidated results for all hypotheses developed in the study, reporting for each relationship the hypothesized direction, the effect type, the estimated standardized path coefficient (β), the t-statistic, the p-value, and the final support decision. Significance is assessed at the operative 5% level, corresponding to a critical t-value of 1.96 for a two-tailed test and a p-value threshold of 0.05, in accordance with the BCa bootstrapping procedure established in Section 5.1 (Hair et al., 2017). The table encompasses the complete set of hypotheses proposed in Chapter 3, comprising total effect hypotheses (H1a, H1b), S→O and O→R pathway hypotheses (H2a, H2b, H3a, H3b), and moderation hypotheses (H4a, H4b, H4c, H4d). Moreover, the table also presents an assessment of the control variables Luxury Interest and Visual Appeal, providing a structured empirical evaluation of the S-O-R framework prior to the broader theoretical and managerial interpretation of the findings in Chapter 6.

Table 5.20. Direction, β , t -statistic, p -values, and hypothesis decision (Source: SmartPLS 4).

Hypothesis	Dir.	Effect Type	β	T stat.	P-val.	Decision
H1a: AI_Disclosure -> WTP	-	Total effect	-0.127	0.778	0.437	Not Supported
H1b: AI_Disclosure -> BRAND_AUTH	-	Total effect	-0.549	3.958	0.000	Supported
H2a: AI_Disclosure -> PER_EFF	-	Direct effect	-0.835	6.640	0.000	Supported
H2b: PER_EFF -> WTP	+	Direct effect	0.162	2.400	0.016	Supported
H2b: PER_EFF -> BRAND_AUTH	+	Direct effect	0.233	2.807	0.005	Supported
H3a: AI_Disclosure -> MOR_DISG	+	Direct effect	0.333	2.400	0.016	Supported
H3b: MOR_DISG -> WTP	-	Direct effect	-0.006	0.057	0.954	Not Supported
H3b: MOR_DISG -> BRAND_AUTH	-	Direct effect	-0.170	2.207	0.027	Supported
H4a: NEED_UNIQ $\times$ AI_Disclosure -> PER_EFF	-	Int. path coeff.	0.023	0.143	0.887	Not Supported
H4b: NEED_UNIQ $\times$ AI_Disclosure -> MOR_DISG	+	Int. path coeff.	-0.110	0.468	0.640	Not Supported
H4c: NEED_UNIQ $\times$ AI_Disclosure -> PER_EFF -> WTP	-	Mod. med. Index	0.004	0.134	0.894	Not Supported
H4c: NEED_UNIQ $\times$ AI_Disclosure -> PER_EFF -> BRAND_AUTH	-	Mod. med. Index	0.005	0.134	0.894	Not Supported
H4d: NEED_UNIQ $\times$ AI_Disclosure -> MOR_DISG -> WTP	-	Mod. med. Index	0.001	0.022	0.982	Not Supported
H4d: NEED_UNIQ $\times$ AI_Disclosure -> MOR_DISG -> BRAND_AUTH	-	Mod. med. Index	0.019	0.391	0.696	Not Supported
CV: LUX_INT -> WTP		Direct effect	0.123	1.753	0.080	Not Supported
CV: LUX_INT -> BRAND_AUTH		Direct effect	0.075	1.036	0.300	Not Supported
CV: VIS_APP -> WTP		Direct effect	-0.049	0.482	0.630	Not Supported
CV: VIS_APP -> BRAND_AUTH		Direct effect	0.225	2.475	0.013	Supported

Note: Dir. = hypothesized direction; Int. path coeff. = interaction path coefficient; Mod. med. index = index of moderated mediation; CV = control variable, to which no directional hypothesis is assigned. Significance assessed via two-tailed test at $p < 0.05$, critical t -value = 1.96. p -values reported as 0.000 indicate $p < 0.001$.

As referenced in Table 5.20, the total effect of AI_Disclosure on WTP is negative and not supported ($\beta = -0.127$, $p = 0.437$), indicating that H1a is not confirmed. Nevertheless, the direction of the total effect remains consistent with the negative prediction. As established in Section 5.3.2, a direction inversion occurs between the positive direct path and the negative total effect, suggesting that the mediating organism-level pathways reverse the sign of the direct stimulus-response relationship. Consequently, as reported in Section 5.3.4, the total indirect effect of AI_Disclosure on WTP is statistically significant, confirming that the negative influence of the stimulus on consumers' economic evaluations operates through the cognitive and affective pathways rather than directly. Conversely, the total effect of AI_Disclosure on BRAND_AUTH is negative and significant ($\beta = -0.549$, $p < 0.001$), supporting H1b and confirming that the disclosure of AI-generated visual content reduces perceived Brand Authenticity relative to the human photography condition.

When analyzing the cognitive pathway, AI_Disclosure significantly and negatively predicts PER_EFF ($\beta = -0.835$, $p < 0.001$), supporting H2a. Furthermore, PER_EFF positively and significantly predicts both WTP ($\beta = 0.162$, $p = 0.016$) and BRAND_AUTH ($\beta = 0.233$, $p = 0.005$), supporting H2b and confirming that higher effort attributions translate into stronger Brand Authenticity evaluations and greater WTP. Accordingly, H2 is supported, confirming that PER_EFF mediates the negative effect of AI_Disclosure on both response constructs. As detailed in Section 5.3.4, the mediation type differs across the two outcomes under the Zhao et al. (2010) framework: the AI_Disclosure $\rightarrow$ PER_EFF $\rightarrow$ WTP chain operates as indirect-only mediation, indicating that the negative economic consequence of AI disclosure depends entirely on the prior cognitive devaluation of Perceived Effort. Instead, the AI_Disclosure $\rightarrow$ PER_EFF $\rightarrow$ BRAND_AUTH chain represents complementary mediation, indicating that PER_EFF partially but not fully accounts for the reduced Brand Authenticity perception, suggesting the potential presence of an additional mediating mechanism, as discussed in Chapter 6.

Analyzing the affective pathway, AI_Disclosure significantly and positively predicts MOR_DISG ($\beta = 0.333$, $p = 0.016$), supporting H3a and confirming that AI disclosure increases consumers' Moral Disgust responses relative to the human photography condition. However, MOR_DISG exhibits divergent effects: it significantly and negatively predicts BRAND_AUTH ($\beta = -0.170$, $p = 0.027$), while its effect on WTP is negligible and not supported ($\beta = -0.006$, $p = 0.954$), making H3b partially supported. Accordingly, H3 is partially supported. This partial support derives from the significance of the individual component paths (specifically H3a and the direct MOR_DISG $\rightarrow$ BRAND_AUTH structural path within H3b) rather than from a statistically confirmed indirect mediation chain, given that both specific indirect effects through MOR_DISG remain non-significant. The AI_Disclosure $\rightarrow$ MOR_DISG $\rightarrow$ WTP chain is classified as no-effect non-mediation under the Zhao et al. (2010) framework, as neither the indirect nor the direct effect of AI_Disclosure on WTP reaches significance for this chain. Conversely, the AI_Disclosure $\rightarrow$ MOR_DISG $\rightarrow$ BRAND_AUTH chain is classified as direct-only non-mediation, confirming that while MOR_DISG exerts a significant influence on BRAND_AUTH, it does not function as a confirmed mediating pathway between AI_Disclosure and Brand Authenticity. The theoretical implications of both classifications are addressed in Chapter 6.

When analyzing the moderating role of NEED_UNIQ, it is possible to note that none of the four interaction paths reaches statistical significance, rendering H4a, H4b, H4c, and H4d not supported and indicating that consumers' Need for Uniqueness does not function as a boundary condition moderating the transmission of AI disclosure through either the cognitive or affective pathway. Furthermore, as observed in Section 5.3.5, the interaction term coefficients demonstrate a directional reversal relative to the hypothesized directions, suggesting that higher levels of Need for Uniqueness appear associated with an attenuated rather than amplified stimulus response. This is a pattern contrary to the theorized sensitivity to the standardization signal carried by AI disclosure within the luxury consumption context. Consequently, H4 is not supported, and the theoretical implications of the reversal are addressed in Chapter 6.

Lastly, when analyzing the control variables, LUX_INT does not significantly predict either WTP ($\beta = 0.123$, $p = 0.080$) or BRAND_AUTH ($\beta = 0.075$, $p = 0.300$), confirming that the observed effects of AI_Disclosure are not attributable to respondents' general engagement with luxury consumption. Notably, the path from LUX_INT to WTP approaches the significance threshold, a marginal relationship discussed in Section 5.3.6. While VIS_APP does not significantly predict WTP ($\beta = -0.049$, $p = 0.630$), it positively and significantly predicts BRAND_AUTH ($\beta = 0.225$, $p = 0.013$), confirming that perceived visual attractiveness independently influences Brand Authenticity evaluations and validating the inclusion of VIS_APP as a control variable.

Regarding the assessment of the model's explanatory power, as referenced in Table 5.11, PER_EFF records an $R^2 = 0.172$ and an $R^2_{adj} = 0.154$, both falling below the weak benchmark threshold established by Hair et al. (2017). Furthermore, MOR_DISG similarly registers below the weak threshold, with an $R^2 = 0.056$ and an $R^2_{adj} = 0.036$. Accordingly, the variance explained for BRAND_AUTH also remains below the weak benchmark, recording an $R^2 = 0.234$ and $R^2_{adj} = 0.207$, despite representing the highest level of explanatory power within the model. Conversely, WTP records the lowest explained variance across all endogenous constructs, with an $R^2 = 0.038$ and an $R^2_{adj} = 0.003$, falling substantially below the weak classification threshold (Hair et al., 2017). Nevertheless, as discussed in Section 5.3.3, these lower R^2 values are consistent with the experimental nature of the present study, whereby a single binary manipulation may account for only a limited portion of the total variance in consumer responses.

Furthermore, as a complementary measure to R^2 , Table 5.21 presents the f^2 effect sizes for the supported hypotheses, whereby values of 0.02, 0.15, and 0.35 indicate small, medium, and large effects respectively (Cohen, 1988; Hair et al., 2017).

Table 5.21. f^2 effect sizes for supported hypotheses (Source: SmartPLS 4).

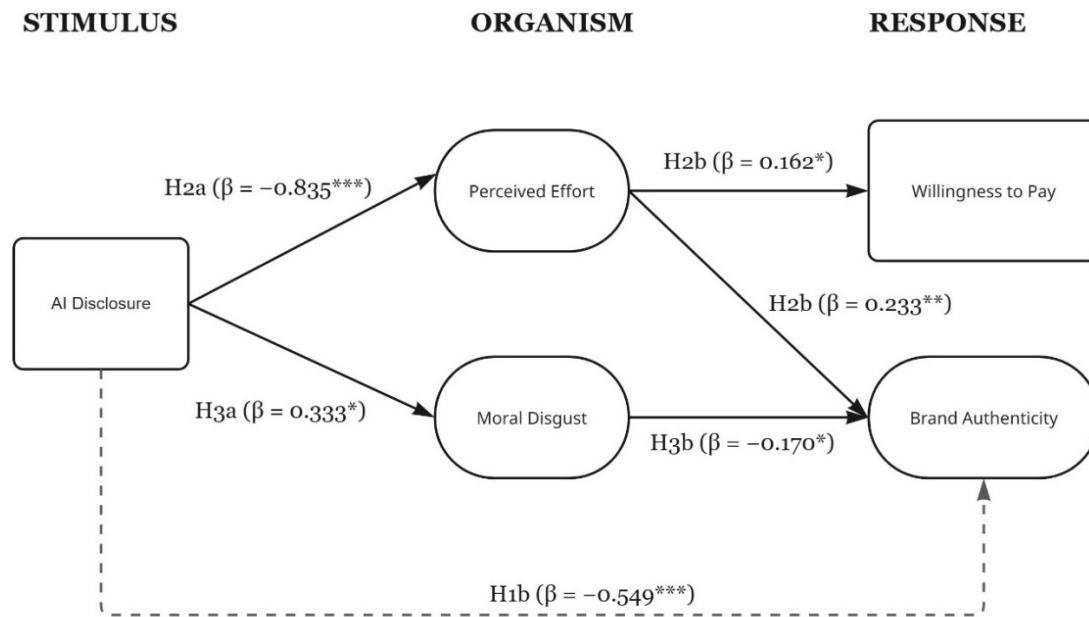
Hypothesis	f^2	Effect
H1b: AI_Disclosure -> BRAND_AUTH	0.022	Small
H2a: AI_Disclosure -> PER_EFF	0.207	Medium
H2b: PER_EFF -> WTP	0.022	Small
H2b: PER_EFF -> BRAND_AUTH	0.057	Small
H3a: AI_Disclosure -> MOR_DISG	0.029	Small
H3b: MOR_DISG -> BRAND_AUTH	0.032	Small

Note: f^2 values drawn from Table 5.12. Effect size classifications follow Cohen (1988) as adopted by Hair et al. (2017): small ($f^2 \geq 0.02$), medium ($f^2 \geq 0.15$), large ($f^2 \geq 0.35$). The f^2 value reported for H1b reflects the direct path contribution of AI_Disclosure to BRAND_AUTH as computed in Table 5.12. The support decision for H1b in Table 5.20 is however based on the total effect ($\beta = -0.549$, $p < 0.001$), which incorporates both direct and indirect contributions and is therefore not directly captured by the f^2 estimate reported here.

As reported in Table 5.21, H2a emerges as the only medium effect within the model ($f^2 = 0.207$), confirming that AI_Disclosure exerts a disproportionately strong individual contribution to the explained variance of PER_EFF relative to all other predictors in the model. The remaining supported hypotheses show small effects, indicating that while each relationship reaches statistical significance, no single predictor beyond H2a accounts for a substantial individual share of variance in its corresponding endogenous construct (Cohen, 1988; Hair et al., 2017). This pattern of small effects across multiple predictors is consistent with the R^2 values reported in Table 5.11, whereby the explanatory power of BRAND_AUTH and WTP reflects the cumulative contribution of several individual effects. Furthermore, the f^2 value associated with H1b reflects only the direct path contribution of AI_Disclosure to BRAND_AUTH rather than the total effect on which its support decision is based. The control variable VIS_APP -> BRAND_AUTH, excluded from Table 5.21, similarly registers as a small effect ($f^2 = 0.057$), suggesting that despite its statistical significance, its individual contribution to the explained variance of BRAND_AUTH remains limited relative to the medium effect observed for the cognitive pathway.

The data analysis conducted throughout Chapter 5 enables the construction of a revised graphical representation of the research model, presented in Figure 5.4, which reflects the graphical structure of the proposed S-O-R framework in Chapter 3.

Figure 5.4. Final graphical representation of the model (Source: the author).



*Note: Control variables (Visual Appeal and Luxury Interest) are omitted from this figure as they do not constitute part of the hypothesized model structure. The significant path from Visual Appeal to Brand Authenticity, established in Section 5.3.6, is accordingly not represented, as its inclusion serves a statistical control purpose within the structural model rather than reflecting a theorized S-O-R relationship. The arrow from Moral Disgust to Brand Authenticity represents a significant direct structural path. As established in Section 5.3.4, this relationship is classified as direct-only non-mediation under the Zhao et al. (2010) framework and does not constitute a confirmed mediating pathway between AI Disclosure and its response constructs. Significance levels: * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$.*

Accordingly, AI Disclosure functions as the exogenous stimulus that initiates the consumer evaluation process, producing a confirmed negative total effect on Brand Authenticity (H1b) and transmitting its psychological and economic consequences through the two parallel mediating constructs of the Organism stage. Within the cognitive pathway (in the Organism), the disclosure of AI-generated visual content is confirmed to reduce consumers' Perceived Effort, such that AI Disclosure negatively predicts Perceived Effort (H2a), which in turn positively predicts both Brand Authenticity and WTP (H2b).

In parallel, the affective pathway operates through Moral Disgust, such that AI Disclosure positively predicts Moral Disgust (H3a), which in turn negatively predicts Brand Authenticity (H3b). Nevertheless, this path represents a significant direct structural relationship rather than a confirmed mediating chain, given that the specific indirect effect through Moral Disgust does not reach statistical significance under the Zhao et al. (2010) framework. Two hypothesis groups do not form part of the empirically supported model structure and are absent from Figure 5.4. The direct total effect of AI Disclosure on WTP (H1a) was not empirically retained, suggesting that the negative influence of AI disclosure on consumers' WTP is not direct in nature but rather operates through the activation of the cognitive and affective mediating pathways. Furthermore, the Need for Uniqueness moderation hypotheses, which theorized that consumers' inherent desire for differentiation and avoidance of similarity would amplify the negative effects of AI disclosure on both the Organism pathways (H4a, H4b, H4c, H4d), were not empirically supported. This pattern suggests that this boundary condition does not meaningfully moderate the transmission of the Stimulus within the present experimental context.

The empirical findings presented throughout Chapter 5 provide evidence for the research questions and theoretical framework established in Chapter 3, enabling a systematic evaluation of the proposed S-O-R model within the GenAI luxury advertising context. Consequently, Chapter 6 interprets these findings against the empirical foundations established in Chapters 2 and 3, from which both theoretical contributions to the GenAI luxury advertising literature and managerial implications for luxury brand communication strategy will be derived.

6. General Discussion

6.1 Overview of Findings

As established in Chapter 3, the regulatory mandate for the disclosure of Artificial Intelligence (AI) introduces a theoretical tension with foundational luxury brand values. The central research objective of the present study is to examine how the mandatory disclosure of AI-generated visual content in luxury advertising shapes consumers' psychological and economic evaluations. This objective was investigated through a Stimulus-Organism-Response (S-O-R) framework, in which the disclosure label functions as the exogenous stimulus, Perceived Effort and Moral Disgust operate as parallel internal mediating constructs, and Willingness to Pay (WTP) and Brand Authenticity constitute the response outcomes, with the moderating role of Need for Uniqueness tested to account for consumer heterogeneity in the transmission of the disclosure effect, as highlighted in Figure 3.1. The empirical evaluation addressed hypotheses tested via Partial Least Squares Structural Equation Modelling (PLS-SEM) on a between-subjects experimental design of $N = 145$, in which the analytical structure evaluated direct effects alongside cognitive and affective mediation before testing moderated mediation (Table 5.20). Taken together, the empirical evaluation yields partial support for the proposed framework. Therefore, the present section synthesizes the principal data analysis from Chapter 5 into discursive form, prior to the theoretical interpretation in Section 6.2.

The partial support of the proposed framework translates into an asymmetric behavior of AI disclosure across the two response constructs. Specifically, the total effect of AI disclosure on Brand Authenticity is confirmed as negative and substantial, which provides empirical support for the argument that explicitly labelling a luxury advertisement as AI-generated substantially undermines consumers' perceptions of brand authenticity and credibility (Table 5.9). This finding contrasts with the non-significant total effect on WTP. Nevertheless, the theoretical relationship cannot be dismissed, given that the coefficient remains consistent with the predicted negative direction even as it does not reach

statistical significance within the present sample. Moreover, the direction inversion between the direct path from AI disclosure to WTP and the negative total effect (Table 5.8; Table 5.9) establishes that the relationship between the stimulus and WTP is not direct in nature. Specifically, the mediating pathways not only transmit but also reverse the direction of the stimulus effect on WTP.

Within this asymmetric pattern, the cognitive mediation chain, operationalized through Perceived Effort, emerges as the most consistently and fully supported component of the empirical model. AI disclosure produces a large negative effect on Perceived Effort, the strongest path coefficient in the entire structural model, confirming that upon receiving the AI disclosure label, consumers systematically attribute lower creative and productive effort to the advertisement's generation (Table 5.9). In turn, Perceived Effort generates significant positive effects on both Brand Authenticity and WTP, yielding a fully confirmed mediation chain across both response constructs, as established in Section 5.3.4 (Table 5.13; Table 5.15). Under the Zhao et al. (2010) mediation classification framework, the mediation is categorized as complementary for Brand Authenticity, where the direct and indirect effects operate in the same negative direction, and as indirect-only for WTP, where the negative indirect contribution of Perceived Effort reverses the direction of the isolated positive direct path. Consequently, this pathway constitutes the only confirmed mechanism through which AI disclosure exerts a statistically significant influence on WTP, establishing Perceived Effort as the primary economic transmission route within the model. This pattern is consistent with the effort heuristic developed in Chapter 2, whereby the cognitive devaluation of the creative process undermines both the psychological legitimacy of the brand's premium positioning and consumers' willingness to pay a financial premium. Furthermore, the total effect of AI disclosure on Brand Authenticity exceeds its isolated direct path, confirming that the mediating pathways amplify rather than merely transmit the stimulus effect on Brand Authenticity.

The affective pathway through Moral Disgust contrasts with the cognitive one, yielding a more complex pattern characterized by partial support and a divergence across the two response constructs. AI disclosure produces a significant positive effect on Moral Disgust (Table 5.9), confirming the stimulus-to-organism link and consistent with the perception

of communicative inauthenticity developed in Chapter 3. Furthermore, Moral Disgust exerts a significant negative effect on Brand Authenticity but a negligible and non-significant effect on WTP. This divergence carries structural consequences for the mediation classification established in Section 5.3.4 (Table 5.13; Table 5.15). Under the Zhao et al. (2010) framework, the affective pathway yields a direct-only non-mediation result for Brand Authenticity: while Moral Disgust operates as a significant direct predictor of Brand Authenticity, the full sequential indirect chain from AI disclosure through Moral Disgust to Brand Authenticity is not confirmed as a mediation pathway. For WTP, neither the direct nor the indirect contribution of Moral Disgust reaches statistical significance, producing a no-effect, no-mediation classification. Taken together, these results establish that the affective pathway does not mediate the stimulus effect on either response outcome, with Moral Disgust instead exerting a selective direct influence on Brand Authenticity and leaving WTP largely unaffected.

Moreover, the present model tested whether Need for Uniqueness amplified AI disclosure effects through both mediating pathways. The moderation does not receive empirical support, given that neither interaction term reaches statistical significance and the moderated indirect effects are absent (Table 5.16; Table 5.17). Furthermore, the main effects of Need for Uniqueness on both Perceived Effort and Moral Disgust are also non-significant (Table 5.9), indicating that the construct does not independently shape either mediating pathway regardless of the disclosure condition. Notably, both interaction coefficients are reversed relative to the theorized predictions in Chapter 3, which suggests that the absence of moderation may reflect a genuine boundary condition on the amplification mechanism proposed in Section 2.5, rather than a limitation of the present sample.

In addition to the moderation boundary condition, Visual Appeal and Luxury Interest were included as control variables to isolate AI disclosure effects from confounding influences. Luxury Interest does not exert a significant effect on either response construct, while Visual Appeal produces a significant positive effect on Brand Authenticity but not on WTP (Table 5.18). The significant effect of Visual Appeal on Brand Authenticity suggests that aesthetic quality contributes independently to perceptions of brand

authenticity in luxury advertising, while the non-significant effect of Luxury Interest on Brand Authenticity suggests that general domain involvement does not systematically bias authenticity evaluations within the present study. Furthermore, the path from Luxury Interest to WTP approaches but does not reach the significance threshold, a finding that warrants examination in future research with larger samples. Taken together, their inclusion ensures that AI disclosure effects on Brand Authenticity are estimated free from confounding influences.

Taken together, these patterns yield partial and asymmetric support for the proposed S-O-R framework. The cognitive pathway emerges as the primary transmission route for the disclosure effect, reaching both response constructs, while the affective pathway operates selectively on Brand Authenticity only and Need for Uniqueness fails to amplify either mediating chain. From this structure, three contributions emerge: the asymmetric dual-pathway structure, the WTP suppression pattern, and the null moderation finding. These constitute the primary objects of theoretical interpretation in Section 6.2 and the basis for the formal contributions to the literature developed in Section 6.4.

6.2 Theoretical Implications

The empirical findings documented in Chapter 5 and synthesized in Section 6.1 yield a pattern of partial and asymmetric support that warrants systematic theoretical interpretation. The asymmetric structure of the dual mediation model, the suppression dynamic underlying the non-significant total effect on WTP, and the absence of Need for Uniqueness moderation each carry distinct implications for the theoretical frameworks developed in Chapters 2 and 3. The discussion that follows interprets these implications in sequence, proceeding from the direct effects through the cognitive and affective mediation blocks to the moderation boundary condition, with reference to the confirmed structural model presented in Figure 5.4 and Table 5.20.

AI disclosure produces asymmetric total effects across the two response constructs, with a substantial negative total effect on Brand Authenticity (H1b: $\beta = -0.549$, $p < 0.001$), confirming that mandatory disclosure of AI-generated visual content systematically lowers consumer perceptions of brand authenticity in the luxury advertising context. This result is consistent with the authenticity penalties documented by Ali et al. (2025) and Brüns & Meißner (2024) following brand adoption of AI-generated content. Conversely, the total effect on WTP does not reach statistical significance (H1a: $\beta = -0.127$, $p = 0.437$), despite maintaining the predicted negative direction (Table 5.9). Moreover, the direct path from AI Disclosure to Brand Authenticity is significant independently of both mediating pathways ($\beta = -0.297$, $p < 0.05$), suggesting that the disclosure label itself functions as a direct signal about the brand's identity-level choices. The brand's decision to deploy automated content generation to replace human creative investment may be read by consumers as a departure from luxury authenticity norms, consistent with Wortel et al.'s (2024) demonstration that persuasion knowledge activation reduces brand attitude through inferences about communicative intent independently of any mediated route.

The cognitive pathway through Perceived Effort provides the most consistent theoretical account of how AI disclosure translates into downstream consequences. The disclosure label suppresses Perceived Effort substantially: the path from AI Disclosure to Perceived Effort (H2a: $\beta = -0.835$, $p < 0.001$, $f^2 = 0.207$) is the single strongest individual coefficient in the model. This suppression is interpretable through the machine heuristic argued by Heimstad et al. (2025), whereby AI source attribution activates a cognitive shortcut that attributes minimal creative investment to algorithmically generated outputs. Moreover, the effect of Perceived Effort on Brand Authenticity (H2b: $\beta = 0.233$, $p = 0.005$, $f^2 = 0.057$) is interpretable through De Kerviler et al.'s (2022) distinction between indexical and iconic authenticity. Perceived effort constitutes the psychological signal through which consumers infer genuine human creative origin, such that when effort attribution is suppressed by the disclosure label, indexical authenticity is removed irrespective of visual quality. Consequently, brand authenticity evaluations deteriorate through the perceived absence of verifiable human authorship rather than through any objective difference in the visual quality of the stimulus. The effect on WTP (H2b: $\beta = 0.162$, $p = 0.016$, $f^2 = 0.022$)

follows a complementary logic: Kruger et al.'s (2004) ambiguity moderator predicts the effort heuristic to be strongest when intrinsic quality is difficult to assess, which characterizes luxury advertising, and Morales' (2005) gratitude-mediated reciprocity mechanism establishes that consumers reward genuinely invested effort with elevated WTP, a reward that disappears when effort is inferred as instrumentally rather than genuinely committed.

The analysis of the affective pathway reveals that AI disclosure produces a significant positive effect on Moral Disgust (H3a: $\beta = 0.333$, $p = 0.016$, $f^2 = 0.029$), replicating the AI-authorship effect of Kirk & Givi (2025) and extending it to a visually mediated luxury advertising domain. The effect size is substantially smaller than that of the cognitive pathway ($f^2 = 0.029$ versus 0.207 for AI Disclosure to Perceived Effort), indicating that the machine heuristic operates with greater force than the moral authenticity violation mechanism in the present context. This contrast may reflect the domain specificity of Kirk & Givi's (2025) findings, which were established in a personal emotional communication context rather than visual advertising. The downstream consequences of Moral Disgust reveal a further theoretical boundary: the path to Brand Authenticity is significant (H3b: $\beta = -0.170$, $p = 0.027$, $f^2 = 0.032$), while the corresponding path to WTP is non-significant (H3b: $\beta = -0.006$, $p = 0.954$). Three converging theoretical accounts explain this pattern: Giner-Sorolla & Chapman (2017) demonstrate that moral disgust responds to evidence of bad moral character and activates avoidance responses whose natural domain is relational and evaluative rather than economic, predicting the pattern obtained. The contagion framework of Amar et al. (2018) offers a complementary explanation: moral disgust triggered by a norm-violating communication spreads via the Law of Contagion to contaminate broader brand authenticity evaluations independently of product quality, accounting for the direct path from Moral Disgust to Brand Authenticity. Under the Zhao et al. (2010) mediation classification framework, the AI Disclosure $\rightarrow$ Moral Disgust $\rightarrow$ Brand Authenticity chain is classified as direct-only non-mediation: the direct path from AI Disclosure to Brand Authenticity is significant, whereas the specific indirect effect through Moral Disgust is not ($\beta = -0.057$, $p = 0.128$). Consequently, Moral Disgust functions as a direct structural antecedent of Brand Authenticity rather than as a confirmed mediating mechanism between the disclosure stimulus and its outcome.

The total effect of AI disclosure on WTP does not reach significance. However, this result does not indicate the absence of economic consequences, as it reflects the simultaneous presence of a significant negative indirect effect and a positive non-significant direct path that partially offset one another at the total effect level. This distribution is reflected in the explained variance of the two response constructs: Brand Authenticity achieves an R^2 of 23.4%, compared to only 3.8% for WTP (Table 5.11), indicating that the structural model accounts for a substantially greater share of variance in Brand Authenticity than in WTP. Within this structure, the indirect-only classification for WTP is particularly relevant from a theoretical standpoint, as it establishes that AI disclosure reduces WTP exclusively through the suppression of Perceived Effort, while the direct path from AI disclosure to WTP is positive and non-significant ($\beta = 0.011$, Table 5.8). This result can be explained by the design features of the study: MERIDIAN carries no prior symbolic equity, and the high-quality AI-generated stimuli may have led some consumers to associate the disclosure with technological sophistication rather than standardization, sustaining an isolated positive direct response on WTP. The negative indirect path, however, can be explained by Morales' (2005) gratitude-mediated reciprocity mechanism: AI disclosure signals automated, cost-minimizing production, suppressing the WTP premium that genuinely invested human effort would generate. Consequently, the non-significant total effect on WTP should not be interpreted as evidence that AI disclosure carries no economic consequences, as the mediation analysis demonstrates that a significant negative indirect effect through Perceived Effort is present within a total effect that appears null only due to the offsetting positive direct path.

The mediation results reveal a consistent difference in how the cognitive and affective pathways transmit the disclosure effect to the two response constructs. The cognitive pathway through Perceived Effort produces confirmed effects on both Brand Authenticity and WTP, while the affective pathway through Moral Disgust produces a confirmed effect on Brand Authenticity only, and the two mechanisms do not lead to the same consumer outcomes despite responding to the same stimulus. Accordingly, the cognitive pathway reaches both outcomes because perceived creative investment shapes not only how consumers evaluate a brand's authenticity but also how much they are willing to pay for it. Kruger et al. (2004) establish that the effort heuristic is stronger when intrinsic quality

is difficult to assess, which characterizes luxury advertising, and Morales' (2005) reciprocity mechanism establishes that consumers translate perceived effort directly into price premiums. The affective pathway, by contrast, reaches only Brand Authenticity because moral disgust motivates consumers to lower their evaluation of the brand rather than to revise their pricing judgments, consistent with Giner-Sorolla & Chapman (2017) and Romani et al. (2012), who each establish that disgust-driven responses operate in the evaluative domain rather than the economic one. This difference also clarifies why the parallel mediator specification is appropriate and why the sequential specification employed by Kirk & Givi (2025) was not applicable in the present study. In their model, perceived authenticity precedes moral disgust as a mediator; in the present study, Brand Authenticity is operationalized as a Response-stage outcome rather than an Organism-stage internal state, a structural difference that makes the sequential specification inappropriate for the present design. The same difference extends to the Visual Appeal control variable, whose significant positive effect on Brand Authenticity ($\beta = 0.225$, $p = 0.013$, $f^2 = 0.057$) but non-significant effect on WTP ($\beta = -0.049$, $p = 0.630$) replicates the relational-evaluative versus economic-pricing pattern at the stimulus level, consistent with To et al.'s (2025) finding that the effort penalty induced by AI disclosure operates independently of objective visual quality. Similarly, Luxury Interest exerts non-significant effects on both Brand Authenticity ($\beta = 0.075$, $p = 0.300$) and WTP ($\beta = 0.123$, $p = 0.080$), although the latter approaches but does not reach the operative significance threshold, a result that may reflect the between-subjects design rather than indicating that luxury orientation plays no role in disclosure responses.

The moderation analysis yields no support for any of the four hypotheses. Neither interaction term reaches statistical significance for Perceived Effort ($\beta = 0.023$, $p = 0.887$) or Moral Disgust ($\beta = -0.110$, $p = 0.640$). The Need for Uniqueness main effects are non-significant on both mediators, and both interaction coefficients are directionally reversed relative to the theorized amplification in Chapter 3. The index of moderated mediation confirms the same pattern across all four conditional indirect effect chains, none of which reaches significance, making H4c and H4d also unsupported (Table 5.17). This pattern is interpretable through three contextual conditions that collectively account for the absence of amplification. First, the disclosure is mandatory and retrospective rather than

a prospective production-mode choice, removing the behavioral channel through which Need for Uniqueness-driven avoidance of similarity operates, a condition absent in Granulo et al. (2021) where participants actively chose between human and robotic production modes. Second, the AI in the present study generates mass-audience advertising rather than personalized consumer-directed output, removing the identity-directed threat that is central to the avoidance mechanism in Loureiro et al.'s (2024) framework. Third, MERIDIAN carries no established brand equity and consumers held no prior identity investment in it. The amplification logic requires precisely such prior investment: high-need-for-uniqueness consumers extend their self-concept through established luxury brand associations, and it is the perceived erosion of those associations that activates the avoidance mechanism. Without a brand in which consumers had already invested symbolically, this activation condition was absent (Belk, 1988; Snyder & Fromkin, 1977). Conversely, the directional reversal of both interaction coefficients appears informative: in the context of a novel brand with high-quality AI-generated stimuli and no prior homogenization signal, high-need-for-uniqueness consumers may have read the disclosure as a marker of technological distinctiveness rather than a standardization threat, consistent with Sands et al.'s (2022) findings on high-need-for-uniqueness evaluative responses reversing when AI output is perceived as genuinely differentiated. Nevertheless, moderation effects require larger samples to detect reliably, and within the present sample of $N = 145$ the statistical power available to detect small interaction effects was limited.

Taken together, these theoretical findings yield three observations that advance the understanding of AI disclosure effects in luxury advertising beyond what prior contributions have established, specifically studies that examine either a single mediating pathway without testing parallel cognitive and affective mechanisms simultaneously, or that assess economic outcomes through total effects without decomposing the underlying indirect and direct forces. The first finding concerns the asymmetric downstream consequences of the two Organism-stage mechanisms: while both the cognitive pathway through Perceived Effort and the affective pathway through Moral Disgust respond to the

same disclosure stimulus, only the cognitive pathway reaches the economic-pricing outcome, establishing that the two mechanisms are not interchangeable within the S-O-R architecture. This result indicates that future dual-pathway models in GenAI disclosure research should not assume that cognitive and affective sub-pathways lead to the same consumer outcomes simply because they are activated by the same stimulus. The second concerns the economic consequences of AI disclosure. The non-significant total effect on WTP, classified as indirect-only mediation, reflects the partial offsetting of a significant negative indirect effect through Perceived Effort by a small positive direct path rather than the absence of economic consequences. This pattern indicates that the economic penalty of AI disclosure was present but operated exclusively through the suppression of Perceived Effort, remaining undetectable at the total effect level due to the countervailing positive direct path. Morales' (2005) gratitude-mediated reciprocity mechanism accounts for this indirect penalty, with Jin & Tao (2025) providing evidence that generative authorship framing reduces WTP through the suppression of perceived creative investment. The third concerns the boundary conditions of Need for Uniqueness moderation. The systematic null result across all four moderation hypotheses, combined with the directional reversal of both interaction coefficients, suggests that the amplification mechanism proposed by Granulo et al. (2021) and Loureiro et al. (2024) does not operate unconditionally. Three contextual features appear necessary for the mechanism to activate. First, the disclosure must be prospective rather than mandatory and retrospective, so that consumers retain a behavioral choice through which their uniqueness motives can be expressed. Second, the stimulus must be directed at the individual consumer rather than a mass audience, generating the identity-directed threat that Need for Uniqueness-driven avoidance requires. Third, the brand must carry prior symbolic equity in which consumers have already invested, activating the extended self-concept mechanism that the amplification logic depends upon. The directional reversal further suggests that, in the absence of these conditions, high-need-for-uniqueness consumers may respond positively rather than negatively to AI disclosure, a pattern consistent with Sands et al. (2022). Accordingly, the managerial implications arising from these findings are developed in Section 6.3, while the study's formal contribution to the literature is addressed in Section 6.4.

6.3 Managerial Implications

The findings documented in the present chapter establish that the mandatory AI disclosure carries measurable consequences for brand authenticity and willingness to pay in luxury advertising contexts, operating through cognitive and affective mechanisms whose effects reach the two response constructs asymmetrically. These consequences are not uniform across the luxury brand landscape and are shaped by conditions relating to brand equity, value architecture, content personalization, and the framing of the disclosure itself. The findings therefore carry implications not only for how luxury brands answer to the regulatory environment but for how they understand the psychological processes the disclosure decision activates. Therefore, the present section translates the findings into actionable guidance for luxury brand practitioners operating under the mandatory disclosure requirements established by the EU AI Act.

As established in Chapter 2, the EU AI Act's transparency provisions changed disclosure from a voluntary communicative choice into a mandatory and legally enforceable condition for brands deploying GenAI content, with General-Purpose AI (GPAI) compliance obligations effective from August 2025. While the AI Act specifies that disclosure must occur, it does not prescribe how it must be framed or contextualized, which leaves ample room in compliance execution. Therefore, the present study's experimental stimulus was designed to replicate the label format, such that the psychological penalties documented reflect consumers' responses to the most common mandated disclosure format, making the recommendations that follow directly applicable to the conditions luxury brands face under the AI Act. While Wortel et al. (2024) demonstrated in an Instagram advertising context that AI disclosure reduces advertising attitude directly and brand attitude among consumers characterized by high AI aversion, the present study extends this evidence to a mandatory visual luxury advertising context, confirming that disclosure-driven evaluative penalties operate even when the label is externally imposed and the consumer has no prior brand equity on which to anchor their evaluations. This represents a more adverse regulatory scenario than previously tested, and one that is analogous to the conditions luxury brands will face under the EU AI Act.

The first managerial implication concerns the communication strategy that should accompany the mandatory label. This recommendation is motivated not only by regulatory practice but also by the study's empirical finding: brand authenticity is the response construct most affected by mandatory AI disclosure, confirming that the primary managerial risk luxury brands face is reputational and relational rather than immediately economic. The dominant mechanism is cognitive: the disclosure label activates the machine heuristic, suppressing perceived creative effort and removing the indexical authenticity signal through which consumers infer genuine human creative origin (De Kerviler et al., 2022; Heimstad et al., 2025). The secondary mechanism is affective: the disclosure triggers moral disgust as a response to perceived norm violation in a context defined by human craftsmanship norms (a direct structural antecedent of brand authenticity), reducing authenticity evaluations independently of the effort mechanism and without representing a confirmed mediation path between the stimulus and the outcome (Giner-Sorolla & Chapman, 2017; Kirk & Givi, 2025). Consequently, luxury brands should develop a creative process transparency strategy that makes visible the human authorship embedded within AI-assisted production, including the creative brief, the aesthetic direction, the analysis of AI outputs, and the final human approval. Nevertheless, the confirmed direct path from AI disclosure to Brand Authenticity, which operates independently of both mediating mechanisms, implies that the combined communication strategies are unlikely to eliminate the brand authenticity penalty in its entirety. Luxury brands should therefore treat the strategies described as partial mitigations of the disclosure effect rather than as complete solutions.

While the cognitive pathway warrants primary managerial attention given its greater effect magnitude, the affective mechanism also warrants a dedicated communicative response. Since moral disgust arises when consumers perceive that a brand has substituted human authorship with automated production, luxury brands can partially mitigate this affective response by proactively communicating that their AI deployment operates as a tool that assists rather than displaces human creative judgment (Hermann & Puntoni, 2025). This communication allows the mandatory label to coexist with a narrative of creative agency, reducing the perceived violation that Kirk & Givi (2025) identify as the initiating condition for the disgust response. Furthermore, evidence from

To et al. (2025) suggests that the perceived effort penalty associated with AI disclosure can be partially attenuated by elevating the creative ambition of the disclosed content itself. Specifically, their third study demonstrates that when luxury brands use AI to generate highly creative rather than standard-creative advertising imagery, the negative impact on perceived effort and advertising authenticity is significantly reduced. This result implies that creative distinctiveness in GenAI content operates as a complementary mitigation strategy alongside transparency, as highly original imagery partially offsets the machine heuristic by signaling genuine aesthetic investment.

The second implication concerns the role of visual quality in luxury advertising under mandatory disclosure conditions. As established in Section 6.2, Visual Appeal exerts a significant positive effect on Brand Authenticity but a non-significant effect on WTP, confirming that aesthetic quality contributes independently to authenticity evaluations while leaving economic valuations unaffected. This pattern is consistent with De Kerviler et al.'s (2022) indexical and iconic authenticity distinction and with Park & Ahn's (2024) finding that while AI-generated luxury imagery can effectively engage consumer-brand relationships through self-brand connections, human-authored communications retain superior explanatory power. High visual quality alone is insufficient to sustain brand authenticity when the indexical signal of human creative investment is removed by the disclosure label. Accordingly, luxury brands should not treat continued investment in visual quality as a substitute for the process transparency strategy described, as evidence establishes that aesthetic quality and perceived human creative origin constitute independent predictors of Brand Authenticity.

The third implication concerns technology investment priorities for luxury brands operating under mandatory disclosure. The preceding discussion may raise a question about whether investing in higher-quality GenAI tools might attenuate the penalties that mandatory disclosure produces. The present study demonstrates that it does not. As established in Chapter 4, the study employed FLUX.2 [dev], an advanced open-source image generation model available at the time of writing with a carefully engineered

prompt, producing a photorealistic stimulus whose visual quality is consistent with the standards that Hartmann et al. (2025) document for state-of-the-art AI-generated imagery. Nevertheless, the disclosure label produced significant and consistent penalties on Brand Authenticity and WTP pathways regardless of image quality parity. The persistence of the disclosure penalty despite the high photorealism of the FLUX.2 [dev] stimulus, consistent with the AI-human quality parity documented by Hartmann et al. (2025), establishes that consumer responses to AI disclosure operate independently of the AI model employed. This result converges with To et al.'s (2025) finding that the effort penalty persists under quality parity conditions, and extends it by demonstrating that the same pattern holds even when a more capable AI image generation model is used, suggesting that the mechanism is robust to improvements in the technical quality of the model deployed. This pattern is consistent with Hagtvedt & Patrick's (2008) observation that luxury aesthetic evaluations are determined by inferred human skill and creative origin, such that authentic authorship perceptions are not recoverable through improvements in output quality. Consequently, luxury brands cannot expect that investment in more capable GenAI models will eliminate the disclosure penalties documented here, as the mechanism is activated by the AI label itself and not by any quality deficiency in the output. This reframes the investment question that mandatory disclosure presents for luxury brands: the primary strategic expenditure does not lie in the continuous improvement of GenAI models, but in the communication strategy surrounding the mandatory label. Addressing the cognitive pathway, this encompasses the documentation and communication of the human creative process embedded within AI-assisted production, including the design brief, the prompt engineering, the evaluation of generated outputs, and the final selection. Addressing the affective pathway, it concerns the proactive framing of the brand's AI adoption philosophy in a manner that positions the technology as operating under human creative authority (Hermann & Puntoni, 2025). These investments target the two confirmed mechanisms through which the disclosure label suppresses brand authenticity evaluations, and represent the most direct mitigation strategy available within the EU AI Act mandatory disclosure.

The fourth implication concerns the commercial significance of the WTP pathway findings and their implications for luxury pricing strategy. The mediation analysis reveals a genuine economic penalty operating through the cognitive pathway, partially masked by a countervailing positive direct path. The distributional WTP analysis in Section 5.4 established that the mean WTP in the human condition (€435.60) exceeds the GenAI condition (€398.63) by approximately €37, with both values falling substantially below comparable luxury market reference prices, including Gucci (€1,650), Louis Vuitton (€2,150), and Prada (€3,100). This divergence may in part be explained by the absence of brand equity in the fictitious brand design. These reference prices illustrate the premium pricing structure that the indirect WTP penalty documented here may progressively undermine for established luxury brands, as the suppression of perceived creative effort reduces the psychological justification consumers require to sustain high price commitments. To the best of the author's knowledge, no prior AI disclosure study has positioned its WTP findings against comparable real-market pricing data. The present study's analysis therefore provides a more commercial account of the economic risk involved. However, for established luxury brands, the indirect WTP penalty may carry commercial consequence. Simon & Fassnacht (2019) establish the premium price architecture as constitutive of the luxury value proposition, such that any mechanism suppressing the perceived justification for a price premium carries commercial consequences at scale. This risk is further supported by Belanche et al. (2025), who demonstrate that AI disclosure produces diminished commercial outcomes in high-involvement and hedonic purchasing contexts, conditions that characterize luxury advertising consumption. Furthermore, it is reasonable to expect that the positive direct path may diminish over time as AI disclosure becomes routine under the EU AI Act, leaving the negative indirect penalty more visible at the total level. Consequently, luxury brands should treat the indirect WTP penalty as a relevant risk requiring constant active monitoring. The urgency of this recommendation is reinforced by a temporal consideration. The positive direct path from AI disclosure to WTP is interpretable as a novelty effect, reflecting consumers' partial association of AI disclosure with technological innovation in a context where mandatory labelling is new. As disclosure becomes routine under the EU AI Act, this novelty response is likely to attenuate, leaving

the negative indirect effect through Perceived Effort more visible at the total level. Since this compensating effect is expected to diminish as mandatory disclosure becomes routine, luxury brands operating in this early phase of the regulatory cycle retain a greater incentive to invest in the communication strategy described before the negative indirect effect becomes fully visible at the total level.

The fifth implication concerns the differentiation of AI adoption strategy across brand equity levels and production stages. The suggested monitoring, however, does not apply uniformly across the luxury brand landscape, given that the disclosure risk appears to vary with brand equity level and value architecture. The penalty is expected to be more severe for established luxury brands with accumulated heritage and consumer identity investment, where consumers have embedded established brand associations within their self-concept, combining authenticity violations with identity-level disruptions (Belk, 1988; Snyder & Fromkin, 1977). Prior AI disclosure research has not differentiated risk profiles across brand equity levels within a single experimental framework. Accordingly, the present study's fictitious brand design, while representing a methodological constraint, provides a case against which established brand penalty severity can be theoretically projected. Conversely, for new or emerging luxury brands, the baseline authenticity expectation has not yet been formed, and the penalty may be partially attenuated. This dynamic extends Xu & Mehta's (2022) finding that the AI penalty is most pronounced in emotionally grounded luxury categories by specifying the dual-pathway mechanism through which this penalty operates in a mandatory advertising disclosure context rather than a product design context. The null moderation result further refines both Granulo et al.'s (2021) and Loureiro et al.'s (2024) Need for Uniqueness amplification arguments. While Granulo et al. (2021) demonstrated Need for Uniqueness sensitivity through prospective production-mode choice, Loureiro et al. (2024) demonstrated it through personalized AI consumer experiences. The present study identifies three conditions, prospective disclosure, identity-directed content, and prior brand equity, whose joint absence prevents the amplification mechanism from activating. This null result carries an additional positive practical implication. Since the disclosure penalty

does not appear to intensify systematically among high-need-for-uniqueness consumers within the present context, luxury brands are not required to develop segment-specific communication strategies for uniqueness-motivated audiences. The transparency and framing approach described can be deployed uniformly across consumer segments without further differentiation.

Accordingly, luxury brands should differentiate their AI adoption decisions based on whether the application involves consumer-visible outputs that trigger the Article 50 disclosure obligation. AI deployed within internal production processes that do not produce consumer-visible content does not activate the mechanisms identified in the present study and therefore does not generate the brand authenticity and WTP penalties documented here. This approach is consistent with Ali et al.'s (2025) hybrid deployment recommendation developed in a restaurant hospitality context, and with Hermann & Puntoni's (2025) human replacement versus human enhancement framework. The present study extends these recommendations to the luxury advertising context by establishing why the distinction between consumer-visible and non-visible AI applications matters strategically. Since both the cognitive and affective penalties documented are activated by the disclosure label rather than by any quality deficiency in the generated image, AI adoption directed toward production stages that do not produce consumer-visible outputs avoids both mechanisms entirely. For established luxury brands with emotionally grounded value architectures, this represents the lowest-risk approach to AI adoption under the current regulatory framework. Where consumer-visible AI deployment is strategically preferred and disclosure is legally mandatory, the communication and framing strategies previously described remain directly applicable.

Taken together, the five implications developed in the present section establish that the disclosure label constitutes the primary risk variable for luxury brands deploying GenAI advertising content under the EU AI Act, and that the most direct mitigation strategy lies in the communication infrastructure surrounding it rather than in the AI model employed. Therefore, the following section develops the contributions to theory of the present study.

6.4 Contributions to Theory

The present study sits at the intersection of five theoretical domains represented in Figure 2.2: luxury brand value architecture, GenAI in marketing and advertising, cognitive consumer responses to AI authorship attribution, affective consumer responses to AI authorship attribution, and consumer heterogeneity through the Need for Uniqueness. Three theoretical contributions emerge from the empirical evaluation: an asymmetric dual-pathway structure in which cognitive and affective mechanisms reach the two response constructs differently, a WTP suppression pattern revealing that null total effects may conceal a significant indirect economic penalty, and a boundary condition specification for the Need for Uniqueness amplification mechanism.

Prior cognitive consumer response studies had confirmed the perceived effort mechanism as the primary path through which AI disclosure penalizes luxury advertising and creative output evaluation, but exclusively through single-pathway designs that could not test whether a simultaneous affective mechanism produces equivalent or different downstream consequences (Heimstad et al., 2025; To et al., 2025). Furthermore, Kirk & Givi (2025) had established moral disgust as a response to AI authorship in personal text-based communication rather than in visually mediated luxury advertising, and without a concurrent cognitive pathway. Prior studies examining AI disclosure effects in social media advertising contexts, such as Wortel et al.'s (2024), had similarly operated under voluntary disclosure conditions in which consumer responses partly reflected inferences about the brand's communicative intent, leaving untested whether the same mechanisms persist when disclosure is externally imposed under regulatory obligation. Xu & Mehta (2022) had attributed the AI penalty in emotionally grounded luxury categories to the primacy of emotional value without specifying the dual mechanism through which it operates under mandatory advertising disclosure, and Ali et al. (2025) applied the S-O-R framework to voluntary AI disclosure in a hospitality context with a single mediating construct.

The present study advances both the cognitive and affective response literatures simultaneously by confirming the machine heuristic mechanism under mandatory

disclosure conditions within a parallel mediation model, establishing that the cognitive pathway reaches both response outcomes independently of the brand's disclosure intent. Moreover, it extends moral disgust from personal text-based communication to visually mediated luxury advertising, reclassifying the latter as a direct structural antecedent of Brand Authenticity under Zhao et al.'s (2010) framework (Amar et al., 2018; Giner-Sorolla & Chapman, 2017; Romani et al., 2012). A residual direct path from AI disclosure to Brand Authenticity, operating independently of both mechanisms, further establishes that the disclosure label functions as a direct evaluative signal beyond its mediated consequences. The study also specifies, against Xu & Mehta (2022), the dual mechanism through which the AI penalty operates in mandatory advertising disclosure rather than product design, extending Ali et al.'s (2025) S-O-R application to a mandatory luxury context and establishing that cognitive and affective pathways are not interchangeable across response outcomes. Finally, the confirmation that visual quality parity does not attenuate the disclosure penalty establishes that the mechanism operates at the level of disclosed production origin rather than output quality, consistent with Hagtvedt & Patrick's (2008) finding that luxury aesthetic evaluations depend on inferred human creative origin.

Prior studies examining the economic consequences of AI disclosure had reported either confirmed negative effects on WTP or purchase intention (Belanche et al., 2025; Jin & Tao, 2025; Loureiro et al., 2024) or null total effects that prior designs could not further decompose, such that the absence of a significant total effect could not be distinguished from the presence of a masked indirect penalty. Furthermore, the luxury brand value architecture literature, while establishing premium pricing as contingent on perceived value justification, had not identified the mechanism through which mandatory AI disclosure specifically threatens this justification. The present study advances the GenAI in marketing literature by demonstrating that a non-significant total effect on WTP does not indicate the absence of economic consequences under mandatory disclosure, representing an indirect-only mediation under Zhao et al.'s (2010) framework. Therefore, a significant negative indirect effect through Perceived Effort may be masked by a countervailing positive direct path at the total level. This positive direct path is

interpretable as a novelty effect, suggesting that the suppression pattern may be temporally bounded as mandatory disclosure becomes routine under the EU AI Act. This dynamic is anchored in Morales (2005), who establishes that consumers reward perceived firm effort with higher WTP, providing the basis for inferring that suppressed effort perceptions reduce WTP through the same reciprocity mechanism. The study further advances the luxury brand value architecture literature by establishing Perceived Effort as the only confirmed economic transmission route, with no affective contribution to WTP confirmed, a result consistent with Giner-Sorolla & Chapman (2017) and Romani et al. (2012), who establish that moral disgust operates in the evaluative rather than the economic domain.

Lastly, Granulo et al. (2021) demonstrated that Need for Uniqueness moderates the preference for human over robotic production in symbolic consumption contexts, while Loureiro et al. (2024) demonstrated that it amplifies avoidance of similarity with downstream WTP consequences in AI-enabled personalized experiences. Nevertheless, neither study identified the contextual conditions under which the amplification mechanism activates or fails to activate. The null result across all four Need for Uniqueness interaction terms, combined with the directional reversal of both interaction coefficients, advances the consumer heterogeneity literature by establishing that the amplification mechanism does not operate unconditionally. The study contributes a boundary condition specification identifying three contextual features whose joint absence prevents activation: prospective rather than mandatory retrospective disclosure, identity-directed rather than mass-audience content, and prior brand equity in which consumers have invested. This result constrains the generalizability of prior amplification findings to contexts meeting all three conditions simultaneously. Conversely, the directional reversal of both interaction coefficients, consistent with Sands et al. (2022), introduces the possibility that high-need-for-uniqueness consumers may respond positively rather than negatively to AI disclosure when these conditions are absent. Therefore, the scope and generalizability of these contributions are examined alongside the study's limitations and future research directions in Chapter 7.

7. Limitations and Future Research

The present study's findings are subject to a number of methodological and theoretical boundary conditions that qualify the scope of interpretation. This chapter addresses these conditions across several thematic areas: sample characteristics, experimental design, construct operationalization, and the explanatory scope of the structural model.

Several limitations concern the composition of the present study's participants. A total of $N = 145$ respondents were recruited through non-probability snowball sampling via the researcher's personal social network and consist exclusively of Italian consumers. This recruitment channel tends to produce a younger and more digitally active respondent profile than the general luxury consumer population, which is relevant because prior AI attitude has been demonstrated to amplify the evaluative penalties associated with AI disclosure (Lim & Schmälzle, 2024). Since prior AI attitude was not measured, it is not possible to determine whether its distribution across the treatment group partly drove the observed Perceived Effort and Moral Disgust responses. The respondents' age profile may also contribute independently to the gap between observed WTP values and real luxury market reference prices documented in Section 5.4. Furthermore, the Italian monocultural frame limits cultural generalizability, as luxury consumption motivations and dispositions towards AI in advertising may vary substantially across national contexts. At $N = 145$, the non-significant indirect effect from AI_Disclosure through MOR_DISG to BRAND_AUTH ($\beta = -0.057$, $p = 0.128$) may reflect limited statistical power rather than the genuine absence of the mechanism, though the near-zero indirect effect on WTP through the same pathway ($\beta = -0.002$, $p = 0.960$) is more consistent with genuine absence (Hair et al., 2017). Similarly, the null moderation result across all four Need for Uniqueness interaction terms cannot be conclusively distinguished from insufficient power to detect small conditional effects. Future research employing cross-national probability-based samples under adequately powered designs would allow both

the cultural generalizability of the documented patterns and the formal re-examination of the moderation hypotheses to be addressed simultaneously.

Several features of the experimental design further limit the scope of the present findings. The manipulation was operationalized through a single image of a leather weekender bag for the fictitious brand MERIDIAN, and the study cannot establish whether the documented penalties generalize across product categories with different expectations of human craftsmanship. Categories such as tailoring or jewelry may generate more severe responses, while categories less directly tied to physical craft may yield attenuated effects (Xu & Mehta, 2022). The fictitious brand design was justified to exclude pre-existing brand associations as confounds (Ali et al., 2025; Kirk & Givi, 2025), but the absence of accumulated brand equity means the study cannot capture the identity-based response that operates when consumers are invested in an established brand (Belk, 1988). WTP values in both conditions fall substantially below real luxury market reference prices suggesting the present findings represent a conservative lower-bound estimate of the penalties established brands would face. The cross-sectional design cannot determine whether penalties change over time as consumers adapt to AI disclosure becoming routine, nor empirically confirm the temporal claim developed in Chapter 6 that the positive direct path from AI_Disclosure to WTP reflects a novelty effect rather than a more stable mechanism. Future research tracking consumer responses over time would allow the temporal dynamics of the novelty effect proposed to be directly tested, while studies employing established luxury brand stimuli across multiple product categories would provide a more representative estimate of the documented penalties.

The operationalization of key constructs introduces a further set of limitations. WTP was operationalized as a single open-ended item and the resulting distribution is substantially right-skewed, as documented in Section 5.4. The open-ended format captures hypothetical stated rather than revealed WTP with established consequences for the precision of economic effect estimates, but at the same time the single-item format prevents standard reliability assessment (Schmidt & Bijmolt, 2020). Furthermore, the

Need for Uniqueness moderator was operationalized exclusively through the Avoidance of Similarity sub-dimension of Tian et al.'s (2001) three-dimensional scale, meaning the null moderation result is bounded to this sub-dimension. The Creative Choice Counterconformity and Unpopular Choice Counterconformity dimensions may be more relevant to the identity-directed conditions identified in Chapter 6, given that the amplification mechanism was documented in contexts presupposing active consumer engagement with product origin rather than the passive reception of a mandatory advertising label (Granulo et al., 2021; Loureiro et al., 2024). The selection of Need for Uniqueness as the only individual-difference moderator means the study cannot determine whether individual-difference moderation is absent from this context entirely or operates through a more proximate construct, such as prior AI attitude, technology readiness, or luxury brand involvement. Independently of the confound concern raised in the sample discussion, their absence from the moderation specification also means the estimated paths from AI_Disclosure to Perceived Effort and Moral Disgust may partly reflect the distribution of unmeasured AI dispositions across conditions (Lim & Schmälzle, 2024). Future studies incorporating validated multi-item WTP measures and direct assessments of prior AI attitude may address the primary operationalization constraints identified here.

The study's experimental context introduces further boundary conditions on its conclusions. With only two conditions, the design cannot establish whether presenting the mandatory disclosure differently, for instance by accompanying the required label with contextual information about the human creative process, would reduce the documented penalties. A related constraint concerns the regulatory framing: the disclosure label used in the experiment does not differ in format from the voluntary AI disclosures examined in prior research (Lim & Schmälzle, 2024), meaning the distinction between mandatory and voluntary disclosure cannot be confirmed through direct experimental comparison. This limitation bears directly on the theoretical contribution identified in Chapter 6, which positions the present findings as applicable to the mandatory regulatory context of the EU AI Act. While this framing is supported by the

study's design rationale, it cannot be empirically confirmed without a dedicated experimental contrast. As the stimulus was formatted as a static image, the mechanism through which the label operates may also differ across video content, editorial print, and influencer-generated communications where different engagement dynamics apply. Research directly contrasting mandatory and voluntary disclosure conditions, and extending the analysis across advertising media beyond the static social media format, may address the regulatory scope and medium specificity constraints identified here. The consistency between the present findings and Bui et al.'s (2024) labelled versus non-labelled between-subjects design in a tourism destination context, where disclosure similarly reversed the direction of authenticity evaluations, suggests that the authenticity penalty mechanism documented here may generalize across visual advertising contexts beyond luxury.

Three analytical observations from the structural model define the study's remaining limitations. The model accounts for approximately 3.8% of the variance in WTP, compared to 23.4% for Brand Authenticity, reflecting greater explanatory reach over the relational-evaluative outcome than the economic one. The remaining WTP variance is attributable to factors outside the model's boundaries, including income, prior brand equity associations, and past luxury purchase experience, whose exclusion was a deliberate design choice but limits inferences about the commercial impact of the disclosure penalty. Out-of-sample predictive validity was also not assessed, as the PLSpredict procedure recommended by Hair et al. (2022) requires k-fold holdout partitioning that is impractical at $N = 145$, and the model's explanatory capacity therefore rests exclusively on in-sample R^2 metrics. The direct path from AI_Disclosure to BRAND_AUTH ($\beta = -0.297, p < 0.05$) persists after both mediators are controlled, yielding a complementary mediation classification that indicates the likely presence of at least one additional unspecified mechanism (Zhao et al., 2010). This residual effect is consistent with persuasion knowledge activation (Wortel et al., 2024), but neither persuasion knowledge, perceived manipulation intent, nor self-brand congruence disruption (Ali et al., 2025) were measured, leaving the model incomplete relative to the full set of

mechanisms operating on Brand Authenticity. Finally, the null result across all four Need for Uniqueness interaction terms raises the question of whether the construct was the most appropriate moderator for a mandatory mass-audience disclosure context. The amplification mechanism documented by Granulo et al. (2021) and Loureiro et al. (2024) was established in contexts that differ structurally from the present mandatory mass-audience disclosure setting. Future research incorporating persuasion knowledge as an additional mediating construct and testing alternative individual-difference moderators alongside Need for Uniqueness would provide a more comprehensive account of the mechanisms through which mandatory AI disclosure shapes luxury brand evaluations.

The limitations identified define the boundary conditions of the present study and point towards a coherent research agenda. Cross-national replications with established luxury brands under adequately powered designs, longitudinal experiments tracking consumer adaptation across the EU AI Act's implementation cycle, multi-condition experiments varying disclosure framing and medium, and model specifications incorporating additional candidate mediators and moderators represent the primary directions through which the contributions established in Chapter 6 may be refined and extended. Pursuing these directions carries particular urgency given that the EU AI Act's disclosure obligations for General-Purpose AI systems became enforceable in August 2025, meaning that the consumer response patterns documented are already relevant for luxury brands active across the European market. The present study's findings therefore serve not only as a theoretical foundation for future inquiry but as an empirically grounded reference point for practitioners and policymakers navigating a regulatory landscape that continues to evolve.

8. Conclusions

The rapid integration of Generative Artificial Intelligence (GenAI) into luxury advertising production has intensified a well-documented tension between technological efficiency and the human craftsmanship and authenticity signals upon which luxury brand value depends (To et al., 2025). Against this background, the EU AI Act addresses the disclosure dimension of this tension at the regulatory level, establishing mandatory transparency requirements for AI-generated advertising content under Article 50 and transforming a previously voluntary communicative decision into a legally enforceable obligation for brands operating across EU Member States (European Parliament, 2024a; Hickman et al., 2024). The present study examined how this mandatory disclosure condition shapes consumers' psychological and economic evaluations in the luxury advertising context, addressing an empirical gap at the intersection of AI disclosure research and luxury brand management that prior work had largely left unexamined. To address this objective, the present study adopted the Stimulus-Organism-Response (S-O-R) framework (Mehrabian & Russell, 1974), operationalizing the mandatory AI disclosure label as the external Stimulus, Perceived Effort and Moral Disgust as parallel Organism-stage mediators representing the cognitive and affective pathways through which the stimulus transmits its effects, and Brand Authenticity and Willingness to Pay (WTP) as the two Response-stage outcomes, with Need for Uniqueness as a moderating variable. The framework was tested through a between-subjects online experiment of $N = 145$ Italian consumers using the fictitious luxury brand MERIDIAN across two conditions differentiated exclusively by the disclosure label (human vs. GenAI). The collected data were subsequently analyzed using PLS-SEM to examine the hypothesized structural relationships.

The empirical evaluation yields partial and asymmetric support for the proposed framework. AI disclosure produces a confirmed substantial negative total effect on Brand Authenticity, while the total effect on WTP does not reach statistical significance. This result reflects a suppression pattern rather than a genuine absence of economic impact, as a confirmed negative indirect effect through Perceived Effort is partially offset by a small positive direct path. Within the Organism stage, the cognitive pathway through

Perceived Effort constitutes the dominant transmission mechanism, with AI disclosure producing the strongest individual path coefficient in the structural model. Perceived Effort in turn generates significant positive effects on both response constructs, yielding complementary mediation for Brand Authenticity and indirect-only mediation for WTP under the Zhao et al. (2010) framework. Through the affective pathway, AI disclosure significantly increases Moral Disgust, which in turn significantly predicts Brand Authenticity but exerts no significant effect on WTP. Neither specific indirect effect through Moral Disgust reaches statistical significance. A residual direct path from AI disclosure to Brand Authenticity persists after controlling for both mediators, indicating the presence of at least one additional unspecified mechanism. Need for Uniqueness moderation is absent across all four tested interaction terms, with both interaction coefficients directionally reversed relative to the theorized amplification pattern.

The findings advance the existing literature across three theoretical contributions. The study first demonstrates that consumer responses to mandatory AI disclosure operate through two distinct psychological pathways, one cognitive and one affective, that produce different outcomes across the response constructs. This finding extends prior research that had examined each pathway independently and without testing them together within a mandatory luxury advertising disclosure context (Kirk & Givi, 2025; To et al., 2025). Furthermore, the study demonstrates that a null total effect on WTP does not indicate the absence of economic consequences, as the mediation analysis confirms a significant negative indirect effect through Perceived Effort that is masked at the total level by a positive direct path, interpretable as a novelty response to a newly mandated regulatory condition that may attenuate as disclosure becomes routine (Morales, 2005). Lastly, Need for Uniqueness moderation is not confirmed across any of the four tested interaction terms, with both coefficients directionally reversed relative to the theorized amplification pattern. This result indicates that the mechanism proposed by Granulo et al. (2021) and Loureiro et al. (2024) does not operate unconditionally, and that its activation appears contingent on three contextual features absent from the present design: prospective rather than mandatory retrospective disclosure, identity-directed rather

than mass-audience content, and prior brand equity in which consumers have already invested.

The findings carry practical implications for luxury brand practitioners navigating the EU AI Act's mandatory disclosure requirements. As the experimental manipulation was limited exclusively to the disclosure label, the documented penalties can be attributed to the mandatory labelling condition rather than to any systematic difference in the quality of the advertising stimuli. Consequently, the strategic priority for luxury brands lies in the communication infrastructure accompanying the mandatory label rather than in improving the GenAI models they utilize. Addressing the cognitive pathway requires luxury brands to communicate the human creative decisions embedded within AI-assisted production, including the creative brief, the prompt development process, and the final selection of outputs. Such communication demonstrates that human expertise and creative judgement shaped the advertising content, and may partially offset the evaluative penalty the mandatory disclosure label generates. Addressing the affective pathway requires framing the brand's AI deployment as operating under human creative authority rather than as a substitution for it, reducing the perceived norm violation that activates Moral Disgust in a luxury craftsmanship context (Hermann & Puntoni, 2025; Kirk & Givi, 2025). The structural model's control variable results provide convergent evidence that investment in visual quality alone is insufficient to offset the authenticity penalty: Visual Appeal exerts a significant positive effect on Brand Authenticity but no significant effect on WTP, confirming that aesthetic quality does not substitute for the indexical signal of human creative authorship that the mandatory disclosure label removes. The persistence of a residual direct path from AI disclosure to Brand Authenticity after both mechanisms are controlled implies that these strategies may reduce but cannot fully eliminate the brand authenticity penalty.

The present study is subject to a number of methodological boundary conditions that qualify the scope of its conclusions. The analytical sample of $N = 145$ Italian consumers, recruited through non-probability snowball sampling, limits the cross-cultural generalizability of the asymmetric pathway pattern and the null moderation result. The

use of a fictitious brand means that the documented effects likely underestimate the consumer responses that established luxury brands with accumulated equity would generate under mandatory disclosure. The cross-sectional design further prevents conclusions about how consumer responses may evolve as mandatory AI disclosure becomes more familiar. Chung et al.'s (2025) evidence that AI-generated images produce a U-shaped satiation trajectory under repeated exposure suggests that longitudinal designs may reveal attenuation of the authenticity penalty documented here as consumers adapt to mandatory disclosure as a standard feature of AI-assisted advertising. Furthermore, WTP was operationalized as a single open-ended item with established limitations for the precision of economic effect estimates (Schmidt & Bijmolt, 2020), and Need for Uniqueness was assessed exclusively through the Avoidance of Similarity sub-dimension. Future research should prioritize replications employing established luxury brands across different national and cultural contexts, longitudinal designs capable of capturing how consumer attitudes evolve as mandatory disclosure becomes standard practice, and expanded models incorporating persuasion knowledge as an additional mediator and prior AI attitude as a moderating variable.

Notwithstanding these boundary conditions, the present study's contributions carry immediate relevance given the active regulatory timeline within which the luxury advertising sector currently operates. The EU AI Act's transparency obligations for AI-generated advertising content became enforceable for General-Purpose AI systems in August 2025, with full compliance extending to all covered categories from August 2026 (European Parliament, 2024a; Hickman et al., 2024). The present study establishes that this mandatory disclosure condition activates two asymmetric psychological pathways in luxury advertising consumers, with the cognitive mechanism generating penalties across both evaluative and economic domains and the affective mechanism confined to evaluative outcomes alone. As mandatory disclosure becomes established practice, the novelty response that currently partially offsets the economic penalty may diminish, strengthening the case for the transparency and communication strategies the present study recommends for luxury brand practitioners.

Bibliography

- Aaker, J., & Schmitt, B. (2001). Culture-dependent assimilation and differentiation of the self: Preferences for consumption symbols in the United States and China. *Journal of Cross-Cultural Psychology*, 32(5), 561–576.
<https://doi.org/10.1177/0022022101032005003>
- Ali, L., Ali, F., Abdalla, M. J., & Alotaibi, S. (2025). Beyond the hype: Evaluating the impact of generative AI on brand authenticity, image, and consumer behavior in the restaurant industry. *International Journal of Hospitality Management*, 131.
<https://doi.org/10.1016/j.ijhm.2025.104318>
- Amaldoss, W., & Jain, S. (2005). Conspicuous consumption and sophisticated thinking. *Management Science*, 51, 1449–1466. <https://doi.org/10.1287/mnsc.1050.0399>
- Amar, M., Ariely, D., Carmon, Z., & Yang, H. (2018). How counterfeits infect genuine products: The role of moral disgust. *Journal of Consumer Psychology*, 28(2), 329–343. <https://doi.org/10.1002/jcpy.1036>
- Atkinson, R., & Flint, J. (2001). Accessing hidden and hard-to-reach populations: Snowball research strategies. *Social Research Update*, 33.
- Bearden, W. O., Netemeyer, R. G., & Haws, K. L. (2011). *Handbook of marketing scales: Multi-item measures for marketing and consumer behavior research*. SAGE Publications.
- Belanche, D., Ibáñez-Sánchez, S., Jordán, P., & Matas, S. (2025). Customer reactions to generative AI vs. real images in high-involvement and hedonic services. *International Journal of Information Management*, 85.
<https://doi.org/10.1016/j.ijinfomgt.2025.102954>
- Belk, R. W. (1988). Possessions and the extended self. *Journal of Consumer Research*, 15(2), 139–168. <https://doi.org/10.1086/209154>
- Black Forest Labs. (2025a, November). *FLUX.2*. Black Forest Labs.
<https://bfl.ai/models/flux-2>

- Black Forest Labs. (2025b, November). *FLUX.2 Documentation*. Black Forest Labs.
https://docs.bfl.ai/flux_2/flux2_overview
- Black Forest Labs. (2025c, November). *FLUX.2 Prompting Guide*. Black Forest Labs.
https://docs.bfl.ai/guides/prompting_guide_flux2
- Borji, A. (2023). Qualitative failures of image generation models and their application in detecting deepfakes. *Image and Vision Computing*, 137, 104771.
<https://doi.org/10.1016/j.imavis.2023.104771>
- Bruner, G. C. (2021). *Marketing scales handbook: multi-item measures for consumer insight research. Volume 11*. GCBII Productions, LLC.
- Brüns, J. D., & Meißner, M. (2024). Do you create your content yourself? Using generative artificial intelligence for social media content creation diminishes perceived brand authenticity. *Journal of Retailing and Consumer Services*, 79.
<https://doi.org/10.1016/j.jretconser.2024.103790>
- Bui, H. T., Filimonau, V., & Sezerel, H. (2024). AI-thenticity: Exploring the effect of perceived authenticity of AI-generated visual content on tourist patronage intentions. *Journal of Destination Marketing and Management*, 34.
<https://doi.org/10.1016/j.jdmm.2024.100956>
- Califano, G., & Spence, C. (2024). Assessing the visual appeal of real/AI-generated food images. *Food Quality and Preference*, 116.
<https://doi.org/10.1016/j.foodqual.2024.105149>
- Chapman, H., & Anderson, A. (2013). Things rank and gross in nature: A review and synthesis of moral disgust. *Psychological Bulletin*, 139, 300–327.
<https://doi.org/10.1037/a0030964>
- Chen, Y., Wang, H., Rao Hill, S., & Li, B. (2024). Consumer attitudes toward AI-generated ads: Appeal types, self-efficacy and AI's social role. *Journal of Business Research*, 185.
<https://doi.org/10.1016/j.jbusres.2024.114867>

- Choi, W. J., & Winterich, K. P. (2013). Can brands move in from the outside? How moral identity enhances out-group brand attitudes. *Journal of Marketing*, 77(2), 96–111. <https://doi.org/10.1509/jm.11.0544>
- Chui, M., Hazan, E., Roberts, R., & Singla, A. (2023). *The economic potential of generative AI: The next productivity frontier*. McKinsey & Company. <https://www.mckinsey.com/capabilities/mckinsey-digital/our-insights/the-economic-potential-of-generative-ai-the-next-productivity-frontier>
- Chung, C., Shin, S., & Chung, N. (2025). Satiation of generative AI images. *Annals of Tourism Research*, 115. <https://doi.org/10.1016/j.annals.2025.104054>
- Cillo, P., & Rubera, G. (2025). Generative AI in innovation and marketing processes: A roadmap of research opportunities. *Journal of the Academy of Marketing Science*, 53(3), 684–701. <https://doi.org/10.1007/s11747-024-01044-7>
- Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed.). Routledge. <https://doi.org/10.4324/9780203771587>
- Comfy. (2025a). *Comfy: The most powerful tool for generative AI*. Comfy. <https://www.comfy.org/>
- Comfy. (2025b, November). *ComfyUI Flux.2 Dev Example*. Comfy. <https://docs.comfy.org/tutorials/flux/flux-2-dev>
- De Kerviler, G., Heuvinck, N., & Gentina, E. (2022). “Make an effort and show me the love!” Effects of indexical and iconic authenticity on perceived brand ethicality. *Journal of Business Ethics*, 179(1), 89–110. <https://doi.org/10.1007/s10551-021-04779-3>
- Diallo, M. F., Ben Dahmane Mouelhi, N., Gadekar, M., & Schill, M. (2021). CSR actions, brand value, and willingness to pay a premium price for luxury brands: Does long-term orientation matter? *Journal of Business Ethics*, 169(2), 241–260. <https://doi.org/10.1007/s10551-020-04486-5>

- Dietvorst, B., Simmons, J., & Massey, C. (2014). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144, 114–126. <https://doi.org/10.1037/xge0000033>
- Dion, D., & Borraz, S. (2015). Managing heritage brands: A study of the sacralization of heritage stores in the luxury industry. *Journal of Retailing and Consumer Services*, 22, 77–84. <https://doi.org/10.1016/j.jretconser.2014.09.005>
- Esposito Vinzi, V., Chin, W., Henseler, J., & Wang, H. (2010). *Handbook of partial least squares: Concepts, methods and applications*. (V. Esposito Vinzi, Ed.). Springer.
- European Parliament. (2024a). EU AI Act. European Parliament. Retrieved February 9, 2026, from <https://explorer.artificialintelligenceact.eu/en/>
- European Parliament. (2024b, June 18). EU AI Act: First regulation on artificial intelligence. European Parliament. Retrieved February 9, 2026, from <https://www.europarl.europa.eu/topics/en/article/20230601ST093804/eu-ai-act-first-regulation-on-artificial-intelligence>
- Feng, W., & Song, X. (2026). Consumers' preference for human-generated versus AI-generated artwork: A social identification account. *Journal of Consumer Behaviour*, 25(1), 118–137. <https://doi.org/10.1002/cb.70063>
- Friestad, M., & Wright, P. (1994). The persuasion knowledge model: How people cope with persuasion attempts. *Journal of Consumer Research*, 21(1), 1–31. <https://doi.org/10.1086/209380>
- Fritz, K., Schoenmueller, V., & Bruhn, M. (2017). Authenticity in branding: Exploring antecedents and consequences of brand authenticity. *European Journal of Marketing*, 51(2), 324–348. <https://doi.org/10.1108/EJM-10-2014-0633>
- Fuchs, C., Prandelli, E., & Schreier, M. (2010). The psychological effects of empowerment strategies on consumers' product demand. *Journal of Marketing*, 74(1), 65–79. <https://doi.org/10.1509/jmkg.74.1.65>

- Fuchs, C., Schreier, M., & Van Osselaer, S. M. J. (2015). The handmade effect: What's love got to do with it? *Journal of Marketing*, 79(2), 98–110.
<https://doi.org/10.1509/jm.14.0018>
- Giner-Sorolla, R., & Chapman, H. A. (2017). Beyond purity: Moral disgust toward bad character. *Psychological Science*, 28(1), 80–91.
<https://doi.org/10.1177/0956797616673193>
- Granulo, A., Fuchs, C., & Puntoni, S. (2021). Preference for human (vs. robotic) labor is stronger in symbolic consumption contexts. *Journal of Consumer Psychology*, 31(1), 72–80. <https://doi.org/10.1002/jcpy.1181>
- Grewal, D., Saturnino, C. B., Davenport, T., & Guha, A. (2025). How generative AI is shaping the future of marketing. *Journal of the Academy of Marketing Science*, 53(3), 702–722. <https://doi.org/10.1007/s11747-024-01064-3>
- Gucci. (2025). Medium duffle bag with Web. Gucci. Retrieved May 10, 2026, from https://www.gucci.com/it/en_gb/pr/men/bags-for-men/duffle-bags-for-men/medium-duffle-bag-with-web-p-758664FACK79768
- Hagtvedt, H., & Patrick, V. M. (2008). Art infusion: The influence of visual art on the perception and evaluation of consumer products. *Journal of Marketing Research*, 45(3), 379–389. <https://doi.org/10.1509/jmkr.45.3.379>
- Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2017). *A primer on partial least squares structural equation modeling (PLS-SEM)* (2nd ed.). Sage Publications.
- Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2022). *A Primer on Partial Least Squares Structural Equation Modeling (PLS-SEM)* (3rd ed.). Sage Publishing.
- Hartmann, J., Exner, Y., & Domdey, S. (2025). The power of generative marketing: Can generative AI create superhuman visual marketing content? *International Journal of Research in Marketing*, 42(1), 13–31.
<https://doi.org/10.1016/j.ijresmar.2024.09.002>
- Hayes, A. F. (2015). An index and test of linear moderated mediation. *Multivariate Behavioral Research*, 50(1), 1–22. <https://doi.org/10.1080/00273171.2014.962683>

- Heimstad, S. B., Wien, A. H., & Gaustad, T. (2025). Machine heuristic in algorithm aversion: Perceived creativity and effort of output created by or with artificial intelligence. *Computers in Human Behavior: Artificial Humans*, 5, 100190. <https://doi.org/10.1016/j.chbah.2025.100190>
- Heitmann, M., Jansen, T. P. J., Reisenbichler, M., & Schweidel, D. A. (2025). Picture perfect: Engaging customers with visual generative AI. *Journal of Marketing*. <https://doi.org/10.1177/00222429251356993>
- Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science*, 43(1), 115–135. <https://doi.org/10.1007/s11747-014-0403-8>
- Hermann, E., & Puntoni, S. (2025). Generative AI in marketing and principles for ethical design and deployment. *Journal of Public Policy and Marketing*, 44(3), 332–349. <https://doi.org/10.1177/07439156241309874>
- Hickman, T., Lorenz, S., & Teetzmann, C. (2024, July 16). Long awaited EU AI Act becomes law after publication in the EU's official journal. White & Case. <https://www.whitecase.com/insight-alert/long-awaited-eu-ai-act-becomes-law-after-publication-eus-official-journal>
- Hinterhuber, A., & Liozu, S. (2017). *Innovation in pricing: Contemporary theories and best practices*. (A. Hinterhuber & S. M. Liozu, Eds.; 2nd ed.). Routledge. <https://doi.org/10.4324/9781315184845>
- Holbrook, M. B., & Hirschman, E. C. (1982). The experiential aspects of consumption: Consumer fantasies, feelings, and fun. *Journal of Consumer Research*, 9(2), 132–140. <https://doi.org/10.1086/208906>
- Huang, M., & Rust, R. (2021). A framework for collaborative artificial intelligence in marketing. *Journal of the Academy of Marketing Science*, 49(1), 30–50. <https://doi.org/10.1007/s11747-020-00749-9>

- Huang, M., & Rust, R. (2022). A Framework for Collaborative Artificial Intelligence in Marketing. *Journal of Retailing*, 98(2), 209–223.
<https://doi.org/10.1016/j.jretai.2021.03.001>
- Jin, Z., & Tao, X. (2025). Framing the future: How AI's creation versus generation visuals affects consumer willingness to pay. *Psychology and Marketing*, 42(10), 2496–2508.
<https://doi.org/10.1002/mar.22242>
- Jones, D. F. (1975). A survey technique to measure demand under various pricing strategies. *Journal of Marketing*, 39(3), 75–77.
<https://doi.org/10.1177/002224297503900316>
- Jung, T., Koghut, M., Lee, E., & Kwon, O. (2025). Artificial creativity in luxury advertising: How trust and perceived humanness drive consumer response to AI-generated content. *Journal of Retailing and Consumer Services*, 87.
<https://doi.org/10.1016/j.jretconser.2025.104403>
- Kapferer, J. N. (2016). The challenges of luxury branding. In F. Dall'Olmo Riley, J. Singh, & C. Blankson (Eds.), *The Routledge companion to contemporary brand management* (pp. 473–491). Routledge.
- Kapferer, J. N., & Bastien, V. (2009). The specificity of luxury management: Turning marketing upside down. *Journal of Brand Management*, 16(5), 311–322.
<https://doi.org/10.1057/bm.2008.51>
- Kapferer, J. N., & Laurent, G. (2016). Where do consumers think luxury begins? A study of perceived minimum price for 21 luxury goods in 7 countries. *Journal of Business Research*, 69(1), 332–340. <https://doi.org/10.1016/j.jbusres.2015.08.005>
- Kapferer, J. N., & Valette-Florence, P. (2021). Which consumers believe luxury must be expensive and why? A cross-cultural comparison of motivations. *Journal of Business Research*, 132, 301–313. <https://doi.org/10.1016/j.jbusres.2021.04.003>
- Keller, K. L. (2009). Managing the growth tradeoff: Challenges and opportunities in luxury branding. *Journal of Brand Management*, 16(5), 290–301.
<https://doi.org/10.1057/bm.2008.47>

- Kim, J.-E., Lloyd, S., & Cervellon, M.-C. (2016). Narrative-transportation storylines in luxury brand advertising: Motivating consumer engagement. *Journal of Business Research*, 69(1), 304–313.
<https://doi.org/https://doi.org/10.1016/j.jbusres.2015.08.002>
- Kirk, C. P., & Givi, J. (2025). The AI-authorship effect: Understanding authenticity, moral disgust, and consumer responses to AI-generated marketing communications. *Journal of Business Research*, 186. <https://doi.org/10.1016/j.jbusres.2024.114984>
- Kruger, J., Wirtz, D., Van Boven, L., & Altermatt, T. W. (2004). The effort heuristic. *Journal of Experimental Social Psychology*, 40(1), 91–98. [https://doi.org/10.1016/S0022-1031\(03\)00065-9](https://doi.org/10.1016/S0022-1031(03)00065-9)
- Kumar, V., Kotler, P., Gupta, S., & Rajan, B. (2025). Generative AI in marketing: Promises, perils, and public policy implications. *Journal of Public Policy and Marketing*, 44(3), 309–331. <https://doi.org/10.1177/07439156241286499>
- 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.
<https://doi.org/10.1016/j.jretconser.2023.103690>
- Lerner, J. S., Small, D. A., & Loewenstein, G. (2004). Heart strings and purse strings: Carryover effects of emotions on economic decisions. *Psychological Science*, 15(5), 337–341. <https://doi.org/10.1111/j.0956-7976.2004.00679.x>
- Leung, F. F., Kim, S., & Tse, C. H. (2020). Highlighting effort versus talent in service employee performance: Customer attributions and responses. *Journal of Marketing*, 84(3), 106–121. <https://doi.org/10.1177/0022242920902722>
- 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>
- Longoni, C., Bonezzi, A., & Morewedge, C. K. (2019). Resistance to medical artificial intelligence. *Journal of Consumer Research*, 46(4), 629–650.
<https://doi.org/10.1093/jcr/ucz013>

- Longoni, C., & Cian, L. (2022). Artificial intelligence in utilitarian vs. hedonic contexts: The 'word-of-machine' effect. *Journal of Marketing*, 86(1), 91–108.
<https://doi.org/10.1177/0022242920957347>
- Louis Vuitton. (2025). Louis Vuitton Keepall Bandoulière 45. Louis Vuitton. Retrieved May 10, 2026, from <https://eu.louisvuitton.com/eng-e1/products/keepall-bandouliere-45-monogram-canvas-000688/M41418>
- Loureiro, S. M. C., Jiménez-Barreto, J., Bilro, R. G., & Romero, J. (2024). Me and my AI: Exploring the effects of consumer self-construal and AI-based experience on avoiding similarity and willingness to pay. *Psychology and Marketing*, 41(1), 151–167. <https://doi.org/10.1002/mar.21913>
- Luffarelli, J., Mukesh, M., & Mahmood, A. (2019). Let the logo do the talking: The influence of logo descriptiveness on brand equity. *Journal of Marketing Research*, 56(5), 862–878. <https://doi.org/10.1177/0022243719845000>
- Magni, F., Park, J., & Chao, M. M. (2024). Humans as creativity gatekeepers: Are we biased against AI creativity? *Journal of Business and Psychology*, 39(3), 643–656. <https://doi.org/10.1007/s10869-023-09910-x>
- Mathwick, C., Malhotra, N., & Rigdon, E. (2001). Experiential value: Conceptualization, measurement and application in the catalog and Internet shopping environment. *Journal of Retailing*, 77(1), 39–56. [https://doi.org/10.1016/S0022-4359\(00\)00045-2](https://doi.org/10.1016/S0022-4359(00)00045-2)
- Mehrabian, A., & Russell, J. A. (1974). *An approach to environmental psychology*. The MIT Press.
- Mogaji, E., & Jain, V. (2024). How generative AI is (will) change consumer behaviour: Postulating the potential impact and implications for research, practice, and policy. *Journal of Consumer Behaviour*, 23(5), 2379–2389. <https://doi.org/10.1002/cb.2345>
- Morales, A. C. (2005). Giving firms an 'E' for effort: Consumer responses to high-effort firms. *Journal of Consumer Research*, 31(4), 806–812. <https://doi.org/10.1086/426615>

- Morhart, F., Malär, L., Guèvremont, A., Girardin, F., & Grohmann, B. (2015). Brand authenticity: An integrative framework and measurement scale. *Journal of Consumer Psychology*, 25(2), 200–218. <https://doi.org/10.1016/j.jcps.2014.11.006>
- Mori, M. (1970). The uncanny valley. *Energy*, (7), 33–35.
- Mori, M., MacDorman, K. F., & Kageki, N. (2012). The uncanny valley [from the field]. *IEEE Robotics & Automation Magazine*, 19(2), 98–100. <https://doi.org/10.1109/MRA.2012.2192811>
- Narayanan, P. (2025). Against the green schema: How Gen-AI negatively impacts green influencer posts. *Psychology and Marketing*, 42(4), 970–986. <https://doi.org/10.1002/mar.22159>
- Nie, X., & Huang, J. (2023). Robots make no effort! Service evaluation and consumer mindset. *Journal of Consumer Behaviour*, 22(2), 365–377. <https://doi.org/10.1002/cb.2138>
- Park, J. K., & Ahn, S. (2024). Traditional vs. AI-generated brand personalities: Impact on brand preference and purchase intention. *Journal of Retailing and Consumer Services*, 81. <https://doi.org/10.1016/j.jretconser.2024.104009>
- Park, J., Oh, C., & Kim, H. Y. (2024). AI vs. human-generated content and accounts on Instagram: User preferences, evaluations, and ethical considerations. *Technology in Society*, 79. <https://doi.org/10.1016/j.techsoc.2024.102705>
- Pedro, Y., Friedmann, E., & Loureiro, S. M. C. (2024). High-end fashion as a social phenomenon: Exploring the perceptions of designers and consumers. *Journal of Retailing and Consumer Services*, 79, 103877. <https://doi.org/10.1016/j.jretconser.2024.103877>
- Peres, R., Schreier, M., Schweidel, D., & Sorescu, A. (2023). On ChatGPT and beyond: How generative artificial intelligence may affect research, teaching, and practice. *International Journal of Research in Marketing*, 40(2), 269–275. <https://doi.org/10.1016/j.ijresmar.2023.03.001>

- Prada. (2025). Sacca da viaggio in Re-Nylon e Saffiano. Prada. Retrieved May 10, 2026, from https://www.prada.com/it/it/p/sacca-da-viaggio-in-re-nylon-e-saffiano/2VC013_2DMH_F0002_V_XOO
- Prasanna, A., & Kushwaha, B. P. (2025a). Generative AI in marketing: Foundations, trends, and future research propositions. *Human Behavior and Emerging Technologies*, 2025, Article 5542513. <https://doi.org/10.1155/hbe2/5542513>
- Prasanna, A., & Kushwaha, B. P. (2025b). Transforming marketing landscapes: A systematic literature review of generative AI using the TCCM model framework. *Management Review Quarterly*. <https://doi.org/10.1007/s11301-025-00486-9>
- Puntoni, S., Reczek, R. W., Giesler, M., & Botti, S. (2021). Consumers and artificial intelligence: An experiential perspective. *Journal of Marketing*, 85(1), 131–151. <https://doi.org/10.1177/0022242920953847>
- Romani, S., Grappi, S., & Dalli, D. (2012). Emotions that drive consumers away from brands: Measuring negative emotions toward brands and their behavioral effects. *International Journal of Research in Marketing*, 29(1), 55–67. <https://doi.org/10.1016/j.ijresmar.2011.07.001>
- Rozin, P., Lowery, L., Imada, S., & Haidt, J. (1999). The CAD triad hypothesis: A mapping between three moral emotions (contempt, anger, disgust) and three moral codes (community, autonomy, divinity). *Journal of Personality and Social Psychology*, 76, 574–586. <https://doi.org/10.1037/0022-3514.76.4.574>
- Russell, S., & Giner-Sorolla, R. (2011). Moral anger, but not moral disgust, responds to intentionality. *Emotion*, 11, 233–240. <https://doi.org/10.1037/a0022598>
- Sands, S., Campbell, C. L., Plangger, K., & Ferraro, C. (2022). Unreal influence: Leveraging AI in influencer marketing. *European Journal of Marketing*, 56(6), 1721–1747. <https://doi.org/10.1108/EJM-12-2019-0949>
- Schmidt, J., & Bijmolt, T. H. A. (2020). Accurately measuring willingness to pay for consumer goods: a meta-analysis of the hypothetical bias. In *Journal of the Academy of Marketing Science* (Vol. 48, Number 3, pp. 499–518). Springer. <https://doi.org/10.1007/s11747-019-00666-6>

- Schmidt, J., Steiner, M., Krafft, M., Eckel, N., & Dahl, D. W. (2024). Hitting the bullseye: Accurately measuring willingness to pay for innovations with open and closed direct questions. *International Journal of Research in Marketing*, 41(2), 383–402. <https://doi.org/10.1016/j.ijresmar.2023.12.003>
- Shukla, P., Singh, J., & Banerjee, M. (2015). They are not all same: variations in Asian consumers' value perceptions of luxury brands. *Marketing Letters*, 26(3), 265–278. <https://doi.org/10.1007/s11002-015-9358-x>
- Sundar, S. S., & Kim, J. (2019, May 2). Machine heuristic: When we trust computers more than humans with our personal information. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems* (pp. 1–9). ACM. <https://doi.org/10.1145/3290605.3300768>
- Simon, H., & Fassnacht, M. (2019). *Price management: Strategy, analysis, decision, implementation*. Springer. <https://doi.org/10.1007/978-3-319-99456-7>
- Snyder, C., & Fromkin, H. (1977). Abnormality as a positive characteristic: The development and validation of a scale measuring need for uniqueness. *Journal of Abnormal Psychology*, 86, 518–527. <https://doi.org/10.1037/0021-843X.86.5.518>
- Snyder, H. (2019). Literature review as a research methodology: An overview and guidelines. *Journal of Business Research*, 104, 333–339. <https://doi.org/10.1016/j.jbusres.2019.07.039>
- Sorosrungruang, T., Ameen, N., & Hackley, C. (2024). How real is real enough? Unveiling the diverse power of generative AI-enabled virtual influencers and the dynamics of human responses. *Psychology and Marketing*, 41(12), 3124–3143. <https://doi.org/10.1002/mar.22105>
- State of California. (2025, September 29). Transparency in Frontier Artificial Intelligence Act. California Legislative Information. Retrieved February 9, 2026, from https://leginfo.legislature.ca.gov/faces/billNavClient.xhtml?bill_id=202520260SB53

- Sundar, S. S. (2020). Rise of machine agency: A framework for studying the psychology of human–AI interaction (HAIL). *Journal of Computer-Mediated Communication*, 25(1), 74–88. <https://doi.org/10.1093/jcmc/zmz026>
- Tajfel, H., & Turner, J. C. (2004). The social identity theory of intergroup behavior. In *Political Psychology* (pp. 276–293). Psychology Press. <https://doi.org/10.4324/9780203505984-16>
- Tian, K. T., Bearden, W. O., & Hunter, G. L. (2001). Consumers' need for uniqueness: Scale development and validation. *Journal of Consumer Research*, 28(1), 50–66. <https://doi.org/10.1086/321947>
- To, R. N., Wu, Y. C., Kianian, P., & Zhang, Z. (2025). When AI doesn't sell Prada: Why using AI-generated advertisements backfires for luxury brands. *Journal of Advertising Research*, 65(2), 202–236. <https://doi.org/10.1080/00218499.2025.2454120>
- Wahid, R., Mero, J., & Ritala, P. (2025). Technology-enabled democratization: Impact of generative AI on content marketing agencies. *Industrial Marketing Management*, 131, 1–16. <https://doi.org/10.1016/j.indmarman.2025.09.007>
- Wang, Y. (2022). A conceptual framework of contemporary luxury consumption. *International Journal of Research in Marketing*, 39(3), 788–803. <https://doi.org/10.1016/j.ijresmar.2021.10.010>
- Wei, C., Peng, L., & Zhu, Q. (2025). Which source is more morally negative? The effect of plagiarism from AI versus human on consumer immoral judgments. *Psychology and Marketing*, 42(10), 2685–2697. <https://doi.org/10.1002/mar.70007>
- Wen, Y., & Laporte, S. (2025). Experiential narratives in marketing: A comparison of generative AI and human content. *Journal of Public Policy and Marketing*, 44(3), 392–410. <https://doi.org/10.1177/07439156241297973>
- Wortel, C., Vanwesenbeeck, I., & Tomas, F. (2024). Made with artificial intelligence: The effect of artificial intelligence disclosures in Instagram advertisements on consumer attitudes. *Emerging Media*, 2(3), 547–570. <https://doi.org/10.1177/27523543241292096>

- Wu, Y. L., & Li, E. Y. (2018). Marketing mix, customer value, and customer loyalty in social commerce: A stimulus-organism-response perspective. *Internet Research*, 28(1), 74–104. <https://doi.org/10.1108/IntR-08-2016-0250>
- Xu, L., & Mehta, R. (2022). Technology devalues luxury? Exploring consumer responses to AI-designed luxury products. *Journal of the Academy of Marketing Science*, 50(6), 1135–1152. <https://doi.org/10.1007/s11747-022-00854-x>
- Zaichkowsky, J. (1985). Measuring the involvement construct. *Journal of Consumer Research*, 12, 341–352. <https://doi.org/10.1086/208520>
- Zhao, X., Lynch Jr., 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. Measurement Scales in Italian (Source: the author).

Construct	Item	Sources
Rivelazione AI	Condizione A: "Fotografia di G. Von Hassel"	Jin & Tao (2025)
	Condizione B: "Immagine generata interamente con Intelligenza Artificiale"	Scala Binaria (0 = "Umano"; 1 = "IA")
Interesse per il Lusso	Quanto ti ritieni interessato/a al mondo dei prodotti e degli accessori di lusso?	Zaichkowsky (1985); Bearden et al. (2011) Scala Likert a sette punti (1 = "Per nulla interessato/a," and 7 = "Estremamente Interessato/a")
Attrazione Visiva	1. Il modo in cui il prodotto è mostrato nell'immagine è attraente.	Mathwick et al. (2001); Bearden et al. (2011) Scala Likert a sette punti (1 = "Fortemente in disaccordo," and 7 = "Fortemente d'accordo")
	2. L'immagine è esteticamente piacevole.	
	3. Mi piace l'aspetto visivo di questa immagine.	
Sforzo Percepito	1. Tempo richiesto per creare l'immagine (1 = Pochissimo tempo, 7 = Moltissimo tempo)	De Kerviler et al. (2022); To et al. (2025) Scala differenziale semantica a 7 punti (ancore specifiche per item)
	2. Impegno richiesto per creare l'immagine (1 = Pochissimo impegno, 7 = Moltissimo impegno)	
	3. Difficoltà nel creare l'immagine (1 = Estremamente facile, 7 = Estremamente difficile)	
Disgusto Morale	1. In che misura il modo in cui questa immagine è stata creata ti fa sentire disgustato/a?	Russell & Giner-Sorolla (2011);

	2. In che misura il modo in cui questa immagine è stata creata ti fa sentire repulsione?	Kirk & Givi (2025) Scala Likert a 7 punti (1 = “Per nulla”, 7 = “Moltissimo”)
	3. In che misura il modo in cui questa immagine è stata creata ti fa sentire nauseato/a?	
Autenticità del Brand	1. Quanto ritieni che il marchio sia Autentico?	Luffarelli et al. (2019); Bruner (2021) Scala Likert a 7 punti (1 = “Per nulla”, 7 = “Moltissimo”)
	2. Quanto ritieni che il marchio sia Affidabile?	
	3. Quanto ritieni che il marchio sia Credibile?	
Disponibilità a Pagare	Qual è il prezzo MASSIMO che saresti disposto/a a pagare per acquistare questa borsa MERIDIAN?	Jones (1975); Fuchs et al. (2010) Domanda aperta (disponibilità a pagare in euro)
Bisogno di Unicità (dimensione: Evitamento della Somiglianza)	1. Evito prodotti o marchi che sono già stati acquistati dal consumatore medio.	Tian et al. (2001); Bearden et al. (2011) Scala Likert a sette punti (1 = “Fortemente in disaccordo,” and 7 = “Fortemente d’accordo”)
	2. Quando prodotti o marchi che mi piacciono diventano estremamente popolari, perdo interesse verso di essi.	
	3. Più un prodotto o marchio è diffuso tra la popolazione generale, meno mi interessa acquistarlo.	

Appendix B. *Summary of Attention and Manipulation Checks in Italian (Source: the author).*

Control Measure	Item	Response Option
Check di attenzione	Per favore, seleziona l'opzione "4 - Né d'accordo né in disaccordo" nella lista qui sotto.	Scala Likert a sette punti (1 = "Fortemente in disaccordo," and 7 = "Fortemente d'accordo") (Risposta corretta: Opzione 4)
Check di manipolazione 1	Ricordi come è stata creata la campagna pubblicitaria che hai valutato all'inizio del questionario?	1. Fotografia tradizionale 2. Generata interamente dall'Intelligenza Artificiale (La risposta corretta dipende dalla condizione assegnata)
Check di manipolazione 2	Oltre alla tecnica di realizzazione, ricordi quale categoria di prodotto era raffigurata nell'immagine pubblicitaria?	1. Una borsa da viaggio (Weekender Bag) 2. Un orologio da polso (Risposta corretta: Opzione 1)