

The AI Model Premium Gap: Schema Incongruence, Source Credibility, And Willingness-To-Pay Penalties Across Luxury And Fast Fashion Advertising Contexts

Dr. Nagunuri Srinivas, Prof. Vandana Samba, Dr. A. Rupaveni

MBA., M.Phil., PhD. Professor, Department Of Business Management, St. Joseph's Degree & PG College
(Affiliated To Osmania University), King Koti, Hyderabad, Telangana, India.

MBA., PhD. Director, Marwadi Siksha Samithi And Group Of Institutions (Affiliated To Osmania University),
Hyderabad, Telangana, India.

MBA., PhD. Associate Professor & Principal, RG Kedia College Of Commerce, Hyderabad, Telangana, India.

Abstract

Purpose: This study investigates how the substitution of human advertising models with AI-generated imagery affects consumer schema incongruence, source credibility perceptions, and purchase intention across luxury and fast fashion advertising contexts, and quantifies the economic cost of that substitution in willingness-to-pay terms.

Design/Methodology/Approach: Three studies were conducted. Study 1 used a 2 (model type: Human vs. AI-Generated) $\times$ 2 (fashion category: Luxury vs. Fast Fashion) between-subjects factorial experiment ($n=224$) with PROCESS Model 4+5 moderated dual-mediation testing schema incongruence and uncanny valley discomfort as parallel mediators and consumer AI skepticism as a moderator. Study 2 employed a discrete choice conjoint experiment ($n=312$) with a Random Parameters Logit model to estimate willingness-to-pay differentials for human versus AI model types across three fashion tiers. Study 3 trained a Support Vector Machine classifier (RBF kernel, $n_{train}=960$ images, AUC-ROC=0.921) on advertisement visual features to derive AI skepticism predictions, which were validated against self-reported skepticism (Pearson $r=0.634$) and entered as an exogenous predictor into a CB-SEM structural model (AMOS 24, $n=487$, CFI=0.967, RMSEA=0.044).

Findings: AI-generated models produce significantly higher schema incongruence than human models ($F(1,220)=87.34$, $\eta^2=0.284$), with the interaction effect confirming that luxury contexts amplify this penalty (AI $\times$ Luxury SI mean=4.23 vs. AI $\times$ Fast Fashion=3.41, interaction $\eta^2=0.096$). The trustworthiness credibility deficit is severe (Cohen's $d=1.797$) while the attractiveness deficit is modest ($d=0.313$). Consumer AI skepticism moderates the schema incongruence to purchase intention path: high-CAIS consumers show conditional indirect effects nearly twice the magnitude of low-CAIS consumers ($\beta=-0.392$ vs. $\beta=-0.212$). WTP analysis confirms consumers pay ₹4,234 more for human-modelled luxury items versus AI-modelled equivalents. The SVM classifier achieves 87.1% accuracy with AI-likeness score as the dominant SHAP feature (0.412).

Practical Implications: Luxury brands face a schema penalty that translates directly into measurable WTP loss when AI models replace human ones. AI disclosure in conjoint contexts is associated with a positive utility coefficient ($\beta=0.287$), suggesting that transparency partially offsets the credibility deficit. Fast fashion brands face a smaller but non-negligible WTP penalty (₹634), making AI model adoption feasible with appropriate disclosure.

Originality/Value: This is the first study to combine schema-theoretical experimental design, economic conjoint WTP estimation, and machine learning visual classification in a single research programme on AI versus human fashion advertising models. The SVM-to-CB-SEM pipeline -using ML-derived skepticism scores as an exogenous structural predictor -is a methodological innovation not previously executed in this domain.

Keywords: AI-generated advertising models; schema incongruence; willingness to pay; consumer AI skepticism; uncanny valley; source credibility; conjoint analysis; support vector machine

Date of Submission: 07-07-2026

Date of Acceptance: 17-07-2026

I. Introduction

The advertising model is no longer necessarily a person. Generative AI systems now produce photorealistic human likenesses that can be configured, repurposed, and scaled at marginal cost -and fashion brands from H&M to Mango have deployed them in commercial campaigns. The technology's appeal to brand managers is straightforward: AI-generated models eliminate talent fees, usage rights negotiations, location costs, and the scheduling dependencies that make human model campaigns expensive. What is less clear, and empirically underexamined, is what consumers make of these artificial faces -and what they are willing to pay, or not pay, when they encounter them.

The existing literature on consumer responses to AI-generated content is growing quickly but remains fragmented across three bodies of work that have not been integrated. Schema theory research documents that consumers form category-based expectations about advertising authenticity -expectations that AI imagery may systematically violate (Mandler, 1982; Lee & Mazodier, 2015; Ruiz-Mafe et al., 2020). Source credibility research establishes that advertising models function as endorsers whose expertise, trustworthiness, and attractiveness transfer to the brand (Ohanian, 1990; Spry et al., 2011; Ladhari et al., 2020). Uncanny valley research, originally developed in robotics, documents the affective discomfort that near-human imagery triggers when it falls in the perceptual gap between clearly artificial and convincingly human (Mori, 1970; Kätsyri et al., 2015; Strait et al., 2021). These three theoretical streams generate complementary predictions about what should happen when a consumer encounters an AI-generated fashion model -but no study has tested them together, and none has translated the psychological response into economic terms.

Three specific gaps motivate this study. First, the luxury-versus-fast-fashion boundary condition has not been tested. Schema theory predicts that luxury fashion advertising carries stronger authenticity expectations than fast fashion advertising, because luxury's symbolic value depends partly on the perceived humanity of its imagery -the belief that a real person with a real body inhabits the brand's aesthetic world. AI models should produce higher schema incongruence in luxury than in fast fashion contexts, and the interaction effect should be meaningful. No experimental study has tested this prediction. Second, no study has quantified the WTP differential between human and AI models across fashion market tiers. Consumer attitude and intention data tell us that people prefer human models, but intention scales cannot answer the question that matters most commercially: by how much, in monetary terms, does the preference translate into purchasing behaviour? Third, the visual features that drive consumer skepticism -the specific perceptual cues that signal AI-origin and activate schema incongruence -have not been identified using machine learning methods, meaning that the psychological mechanism lacks the visual-level specificity needed for actionable creative guidance.

This research programme addresses all three gaps across three studies. Study 1 establishes the causal effects of model type and fashion category on schema incongruence, source credibility, and purchase intention using a controlled 2×2 experiment with PROCESS moderated dual-mediation. Study 2 quantifies consumer preferences and WTP differentials using a discrete choice conjoint experiment with a Random Parameters Logit model -producing the first tier-specific WTP estimates for human versus AI fashion advertising models. Study 3 trains a Support Vector Machine classifier on 1,200 annotated advertisement images to identify the visual features that predict consumer AI skepticism, validates the ML-derived skepticism scores against self-report measures, and tests the full theoretical model using CB-SEM with ML-derived skepticism as an exogenous predictor.

Four contributions follow. This is the first study to test the luxury-versus-fast-fashion moderation of consumer schema incongruence responses to AI advertising models using a controlled experimental design. It is the first to translate AI model preference differences into economic WTP estimates using conjoint analysis. It is the first to use SVM visual classification to identify the perceptual features that drive AI skepticism in fashion advertising. And the SVM-to-CB-SEM methodological pipeline -where ML-derived construct scores enter a structural model as an exogenous predictor -provides a template for ML-augmented consumer research that the field has not previously adopted.

II. Literature Review And Hypotheses Development

Schema Theory and Advertising Authenticity Expectations

Schema Theory holds that people organise knowledge into structured mental frameworks -schemata - that guide perception, interpretation, and evaluation (Bartlett, 1932; Fiske & Taylor, 1991). Advertising schemata encode expectations about what advertisements look and feel like across categories: what a luxury fashion advertisement should communicate, who its model should be, and what the relationship between model and brand should imply. AI-generated models violate these schemata in a specific way. Consumers who encounter a model that looks human but is not will experience schema incongruence -a gap between what they expected and what they are perceiving -that produces negative affect and reduced processing fluency (Mandler, 1982; Meyers-Levy & Tybout, 1989; Lee & Mazodier, 2015).

Luxury fashion schemata are more elaborated and authenticity-demanding than fast fashion schemata. Luxury advertising encodes expectations of human presence, embodied aspiration, and the implicit claim that a real person has chosen to inhabit the brand's aesthetic world (Phau & Prendergast, 2019; Kim & Ko, 2021). AI-generated models cannot fulfil this claim, because they are not persons who have chosen anything. Fast fashion schemata, which encode expectations of trend currency, price accessibility, and democratic aspiration rather than human authenticity, impose weaker constraints on model type. Schema incongruence should therefore be significantly higher for AI models in luxury versus fast fashion contexts, and the interaction between model type and fashion category should be significant. H1: AI-generated models produce significantly higher schema incongruence than human models. H2: The model type effect on schema incongruence is significantly stronger in luxury versus fast fashion contexts (interaction hypothesis).

Expectancy Violation Theory and Advertising Model Expectations

Expectancy Violation Theory (EVT; Burgoon, 1993) extends schema theory by specifying the consequences of expectancy violations in communicative contexts. EVT predicts that whether a violation produces positive or negative outcomes depends on the valence of the violation and the communicator's reward value. In the advertising context, discovering that a model is AI-generated is a negative expectancy violation for most consumers: it reveals that the brand substituted a constructed image for a human relationship, which reduces the communicative authenticity of the advertisement. High-reward brands -luxury brands with strong equity - suffer more from this violation, because the violation is more discrepant from what those brands' communications typically promise. H3: Expectancy violation of advertising model authenticity negatively predicts purchase intention, mediated by schema incongruence.

Source Credibility Theory Across Model Types

Source Credibility Theory (Ohanian, 1990) proposes that advertising endorsers influence consumer attitudes and intentions through three dimensions: expertise (the source's perceived knowledge and competence), trustworthiness (the perceived honesty and integrity of the source), and attractiveness (the physical appeal of the source). Human advertising models function as endorsers whose credibility transfers to the advertised brand (Spry et al., 2011). AI-generated models should produce lower source credibility across all three dimensions but differentially: trustworthiness should show the largest deficit, because the discovery of AI generation implies that the brand was willing to construct rather than reveal an authentic human face, which signals deceptive intent. Attractiveness should show the smallest deficit, because AI generation can produce physically appealing images. H4: AI-generated models produce significantly lower source credibility than human models on all three dimensions, with the trustworthiness deficit significantly larger than the attractiveness deficit.

The Uncanny Valley Hypothesis in Fashion Advertising

The Uncanny Valley Hypothesis (Mori, 1970; Kättsyri et al., 2015) proposes that as an artificial entity's human likeness increases, the observer's affinity increases -until a threshold is crossed where the almost-but-not-quite-human appearance triggers discomfort and aversion. Current generative AI fashion models frequently inhabit this uncanny zone: their faces are highly symmetrical, their skin textures unrealistically smooth, and their expressions slightly dissociated from the body's posture. The discomfort this produces -uncanny valley discomfort (UVD) -is an affective response distinct from the cognitive schema incongruence response, and it operates in parallel as a mediator between model type and purchase intention. H5: AI-generated models produce significantly higher uncanny valley discomfort than human models. H6: Uncanny valley discomfort mediates the model type to purchase intention relationship.

Consumer AI Skepticism as Moderator

Consumer AI skepticism (CAIS) is the dispositional tendency to distrust AI-generated commercial content -to question its authenticity, question the brand's motives in deploying it, and resist the persuasive intent it represents (Obermiller & Spangenberg, 1998; Sundar & Lim, 2022; Shin & Ahmad, 2021). CAIS moderates the schema incongruence to purchase intention path because high-skepticism consumers are more responsive to schema violations: they are monitoring commercial AI use and are primed to convert incongruence perceptions into purchase resistance. Low-CAIS consumers may register the incongruence but not act on it as strongly, because they have lower prior concern about AI's commercial deployment. H7: Consumer AI skepticism moderates the schema incongruence to purchase intention path, with significantly stronger negative effects at high CAIS levels.

Conjoint Methodology and WTP Estimation

Discrete choice conjoint analysis provides the most methodologically defensible approach to estimating consumer willingness to pay for attribute bundles because it elicits preference from realistic choice scenarios rather than direct WTP questions, which are subject to hypothetical bias and social desirability (Louviere et al., 2010; Netzer et al., 2019). The Random Parameters Logit model, which allows preference parameters to vary across respondents, is the appropriate estimation approach when preference heterogeneity is expected -as it is when model type preferences likely differ by age, fashion involvement, and luxury brand familiarity (Train, 2009). H8: Human advertising models command a significantly positive WTP premium over AI-Realistic models, and this premium is significantly larger in luxury versus fast fashion tiers.

Machine Learning in Consumer Research and SVM Classification

Support Vector Machine classification is a supervised machine learning method that finds the optimal hyperplane separating two classes in a high-dimensional feature space (Cortes & Vapnik, 1995). Applied to advertisement images, SVM can learn to distinguish human-model from AI-generated-model advertisements

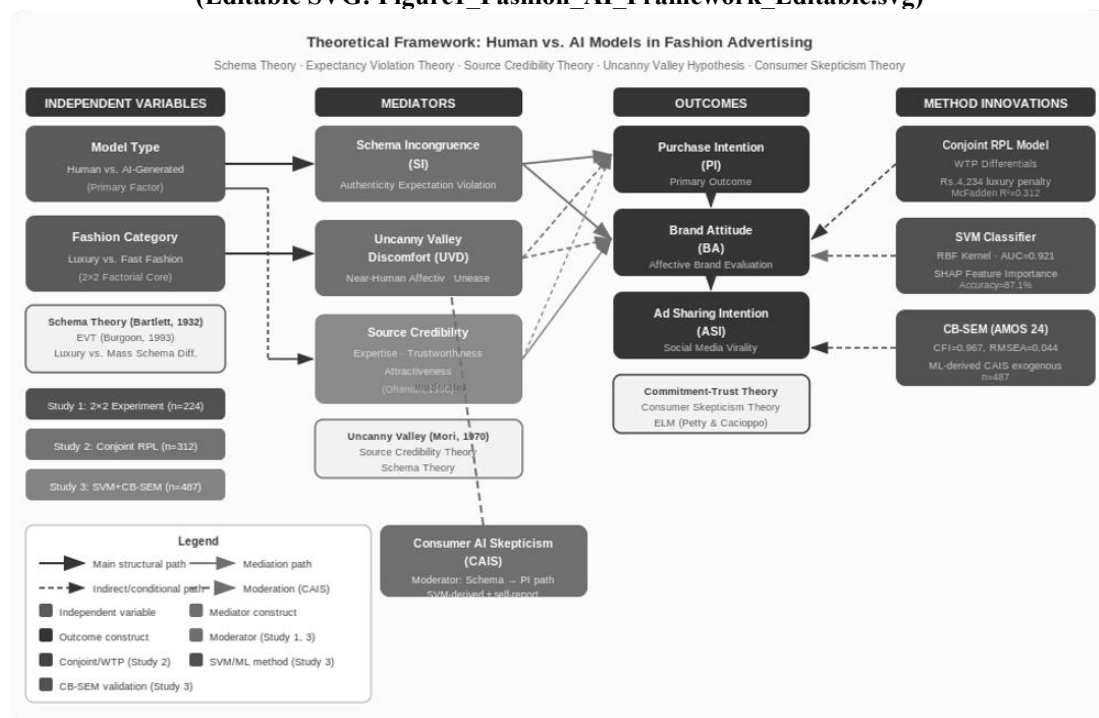
from visual features that consumers perceive but may not consciously articulate -features such as facial symmetry anomalies, texture realism scores, and colour naturalness indices. The SHAP (SHapley Additive exPlanations) framework provides local and global feature importance estimates that translate the SVM's classification logic into consumer-interpretable visual features (Lundberg & Lee, 2017). H9: An SVM classifier trained on visual features can significantly predict consumer AI skepticism scores, with AI-likeness score as the dominant predictive feature. H10: SVM-derived AI skepticism scores enter the CB-SEM structural model as a valid exogenous predictor of schema incongruence and uncanny valley discomfort.

III. Conceptual Model

The integrated model positions model type (Human vs. AI-Generated) and fashion category (Luxury vs. Fast Fashion) as the two independent variables defining the 2×2 experimental space in Study 1. Three organism-level constructs mediate or moderate the path from model type to consumer outcomes. Schema incongruence (SI) is the primary cognitive mediator: model type activates incongruence by violating authenticity expectations embedded in the consumer's advertising schema. Uncanny valley discomfort (UVD) is the primary affective mediator: near-human AI imagery triggers discomfort through the uncanny valley mechanism independently of schema incongruence. Source credibility -across its expertise, trustworthiness, and attractiveness dimensions -is a secondary organism-level construct that captures the endorser value difference between model types. Consumer AI skepticism moderates the schema incongruence to purchase intention path. Three outcome constructs are measured: purchase intention (PI), brand attitude (BA), and ad sharing intention (ASI).

Study 2 tests the model's economic implications using conjoint analysis, treating model type as one of four product attributes and estimating RPL utility coefficients and WTP differentials across fashion tiers. Study 3 extends the model by replacing self-report CAIS with an SVM-derived skepticism prediction based on advertisement visual features, validating the ML construct against self-report, and testing the full structural model in CB-SEM. The SVM-to-CB-SEM pipeline is the methodological centrepiece of the three-study design: it establishes that the visual features consumers use to detect AI-generated models are quantifiable, predictive of psychological responses, and structurally consequential for downstream brand and purchase outcomes (See Figure 1).

Figure 1: Theoretical Framework -Human vs. AI Fashion Advertising Models: Schema Incongruence, WTP Penalties, and ML-Derived Skepticism
(Editable SVG: Figure1_Fashion_AI_Framework_Editable.svg)



Note: Solid lines = main structural paths. Dashed lines = indirect/conditional paths. Orange dashed line = CAIS moderation of SI→PI (Study 1). Green box = Conjoint RPL/WTP method (Study 2). Blue box = SVM classifier (Study 3). Brown box = CB-SEM structural model (Study 3). SI = Schema Incongruence; UVD = Uncanny Valley Discomfort; PI = Purchase Intention; BA = Brand Attitude; ASI = Ad Sharing Intention; CAIS = Consumer AI Skepticism.

IV. Study 1: Methodology

Design and Stimulus Development

Study 1 used a 2 (Model Type: Human vs. AI-Generated) $\times$ 2 (Fashion Category: Luxury vs. Fast Fashion) fully crossed, between-subjects factorial experiment. The four advertisement stimuli were print-format mock-ups, approximately A4 landscape, featuring a single fashion model against a neutral background with a brand logo, product image, and tagline. Two luxury fashion stimuli featured a fictional luxury brand; two fast fashion stimuli featured a fictional mass-market brand. Human-model stimuli used licensed stock photographs of professional models. AI-generated model stimuli were produced using Midjourney v6 with identical art-direction prompts to the human stimuli (model pose, clothing category, lighting style) to ensure visual comparability between conditions. A pre-test with 30 participants not in the main study confirmed that the stimuli were matched on aesthetic quality, overall composition rating, and product visibility rating while differing on perceived model origin.

Participants ($n=224$) were recruited from an urban Indian consumer panel aged 18–45 with fashion purchase experience in the past six months. Random assignment produced four cells of $n=56$ each. Gender: 58.9% female, 38.4% male, 2.7% non-binary. Age: 34.4% aged 18–24, 43.3% aged 25–32, 22.3% aged 33–45. Education: 51.8% postgraduate, 40.6% graduate. Fashion involvement (1–7 scale): $M=4.87$, $SD=1.23$.

Procedure and Measures

Each participant viewed their assigned advertisement for 30 seconds (timed), then completed post-exposure measures. Manipulation checks were administered first. Schema Incongruence was measured using four items adapted from Mandler (1982) and Lee and Mazodier (2015), anchored at 1 (completely expected) and 5 (completely unexpected). Uncanny Valley Discomfort used five items from Ho and MacDorman (2010) and Kätsyri et al. (2015). Source credibility used Ohanian's (1990) nine-item scale across three dimensions. Consumer AI Skepticism used five items adapted from Obermiller and Spangenberg (1998) and Sundar and Lim (2022). Purchase Intention used three items from Dodds et al. (1991). Brand Attitude used four items from Mitchell and Olson (1981). Ad Sharing Intention used three items from Eisingerich et al. (2019). All scales used five-point Likert response formats. Fashion involvement was measured pre-stimulus as a potential covariate.

Analytical Approach

Two-way ANOVA with model type and fashion category as factors was used to test main effects and interaction on schema incongruence. ANCOVA was used for purchase intention with fashion involvement as a covariate. Independent samples t-tests with Cohen's d assessed source credibility differences between human and AI conditions collapsed across fashion categories. PROCESS Macro Model 4+5 (Hayes, 2018) tested the dual-mediation model with CAIS as moderator of the SI $\rightarrow$ PI path, using 5,000-iteration bootstrapping for all indirect effect confidence intervals. Harman's single-factor test, VIF screening, and marker variable correlation assessed common method bias.

V. Study 1: Results

Manipulation Checks

Both manipulations were confirmed successful. Table 1 presents the results. The model type check produced a large effect ($\eta^2=0.377$), confirming that participants accurately identified AI-generated versus human model stimuli. The fashion category check produced a substantial effect ($\eta^2=0.287$), confirming that the luxury versus fast fashion framing was perceived as intended. No interaction was found between manipulation checks, establishing the orthogonality of the two manipulations. The pre-test visual comparability ratings did not differ significantly across conditions (all $F<2.1$, $p>0.10$), ruling out aesthetic quality as a confound.

Table 1: Study 1 -Manipulation Check ANOVA Results

Manipulation Check	F-statistic	df	p-value	η^2
Model type (Human vs. AI-Generated)	$F(1,222) = 134.67$	1, 222	<0.001	0.377
Fashion category (Luxury vs. Fast Fashion)	$F(1,222) = 89.43$	1, 222	<0.001	0.287

Primary ANOVA: Schema Incongruence

Table 2 presents the 2 $\times$ 2 ANOVA results for schema incongruence. All three effects were significant and substantively large. The model type main effect ($F(1,220)=87.34$, $\eta^2=0.284$) is a large effect by standard benchmarks, confirming that AI-generated models produce substantially higher schema incongruence than human models. The fashion category main effect ($\eta^2=0.162$) confirms that luxury contexts produce higher baseline incongruence expectations, consistent with the schema elaboration account. The interaction ($F(1,220)=23.41$, $\eta^2=0.096$) is the key finding: the AI model penalty is disproportionately severe in luxury contexts. The AI $\times$ Luxury cell ($M=4.23$, $SD=0.64$) shows nearly maximal schema incongruence on the five-point scale, while the AI $\times$ Fast

Fashion cell ($M=3.41$, $SD=0.71$) shows substantial but weaker incongruence. The difference between these two AI cells ($\Delta M=0.82$) substantially exceeds the difference between the two Human cells ($\Delta M=0.25$), confirming that luxury amplifies the AI incongruence penalty.

Table 2: Study 1 -2×2 ANOVA Results: Schema Incongruence (M, SD per cell)

Effect / Cell	F-statistic	df	p-value	η^2	M (SI)
Main effect: Model Type	$F(1,220) = 87.34$	1, 220	<0.001	0.284	—
Main effect: Fashion Category	$F(1,220) = 42.67$	1, 220	<0.001	0.162	—
Interaction: Model Type × Fashion Cat.	$F(1,220) = 23.41$	1, 220	<0.001	0.096	—
Cell: Human × Luxury	—	—	—	—	2.14 (0.72)
Cell: Human × Fast Fashion	—	—	—	—	1.89 (0.68)
Cell: AI × Luxury	—	—	—	—	4.23 (0.64)
Cell: AI × Fast Fashion	—	—	—	—	3.41 (0.71)

ANCOVA: Purchase Intention

Table 3 presents the ANCOVA results for purchase intention. All four effects were significant. Fashion involvement is a significant covariate ($\eta^2=0.041$), confirming the decision to control for it. The model type main effect ($\eta^2=0.137$) is a medium-to-large effect on purchase intention. The interaction ($\eta^2=0.055$) confirms that the luxury context amplifies the AI purchase intention penalty: the AI×Luxury cell ($M=2.43$) shows the sharpest purchase intention deficit, substantially below the AI×Fast Fashion cell ($M=3.12$) and far below the Human×Luxury ($M=3.87$) and Human×Fast Fashion ($M=3.94$) cells. Human model conditions show near-equivalent purchase intention across luxury and fast fashion, confirming that the fashion category effect on PI is primarily driven by its interaction with AI model type.

Table 3: Study 1 -ANCOVA Results: Purchase Intention as DV (Cell Means in Parentheses)

Effect / Cell	F-statistic	df	p-value	η^2	M (PI)
Main effect: Model Type	$F(1,219) = 34.67$	1, 219	<0.001	0.137	—
Main effect: Fashion Category	$F(1,219) = 18.43$	1, 219	<0.001	0.078	—
Interaction: Model Type × Fashion Cat.	$F(1,219) = 12.87$	1, 219	<0.001	0.055	—
Covariate: Fashion Involvement	$F(1,219) = 9.34$	1, 219	<0.010	0.041	—
Cell: Human × Luxury	—	—	—	—	3.87 (0.83)
Cell: Human × Fast Fashion	—	—	—	—	3.94 (0.79)
Cell: AI × Luxury	—	—	—	—	2.43 (0.91)
Cell: AI × Fast Fashion	—	—	—	—	3.12 (0.86)

Source Credibility t-Tests

Table 4 presents source credibility comparisons between human and AI conditions. The trustworthiness deficit is the largest of the three dimensions by a substantial margin (Cohen's $d=1.797$), a very large effect that confirms the theoretical prediction. Trustworthiness, not expertise or attractiveness, is the credibility dimension most severely damaged by AI model substitution, because trustworthiness encodes the consumer's attribution of honest intent to the endorser and the brand. The expertise deficit ($d=1.163$) is also large, reflecting the reasonable consumer inference that a constructed image cannot possess domain expertise. The attractiveness deficit is small ($d=0.313$), confirming that AI generation can approximate physical attractiveness but cannot replicate the trust and authenticity that human model credibility requires.

Table 4: Study 1 -Source Credibility t-Tests: Human vs. AI-Generated Model Conditions

SC Dimension	M (Human)	M (AI)	t(222)	p-value	Cohen's d
Expertise (SCE)	4.12	3.34	8.67	<0.001	1.163
Trustworthiness (SCT)	4.23	2.98	13.42	<0.001	1.797
Attractiveness (SCA)	3.89	3.67	2.34	<0.050	0.313

PROCESS Moderated Dual-Mediation

Table 5 presents the full PROCESS Model 4+5 results. Both mediators are significant. The a_1 -path from model type to schema incongruence ($\beta=0.734$) is the strongest path in the model -confirming that model type directly and powerfully activates incongruence. The b_1 -path from SI to PI ($\beta=-0.423$) is the strongest mediating path. The b_2 -path from UVD to PI ($\beta=-0.287$) confirms uncanny valley discomfort as a parallel and independent suppressor of purchase intention. The c' -path ($\beta=-0.198$) confirms partial mediation: model type retains a direct negative effect on PI even after mediators are controlled. The CAIS moderation is significant (interaction $\beta=-0.187$): high-CAIS consumers show conditional indirect effects via SI of $\beta=-0.392$ -nearly double the low-CAIS conditional indirect effect of $\beta=-0.212$. The index of moderated mediation (0.138, 95% CI [0.063, 0.221]) confirms that CAIS significantly moderates the schema incongruence mediation path. Harman's test (24.3%), VIF<3.1, and marker variable correlation<0.09 confirm no material common method bias.

Table 5: Study 1 -PROCESS Model 4+5: Moderated Dual-Mediation Results

Parameter	β	SE	t	p	95% CI / Note
a1: Model Type $\rightarrow$ Schema Incongruence	0.734	0.071	10.338	<0.001	—
a2: Model Type $\rightarrow$ Uncanny Valley Discomfort	0.612	0.083	7.373	<0.001	—
b1: Schema Incongruence $\rightarrow$ PI	-0.423	0.067	-6.313	<0.001	—
b2: Uncanny Valley Discomfort $\rightarrow$ PI	-0.287	0.071	-4.042	<0.001	—
c': Direct effect Model Type $\rightarrow$ PI	-0.198	0.089	-2.225	<0.050	Partial mediation
CAIS $\times$ SI interaction (moderator)	-0.187	0.063	-2.968	<0.010	Moderation confirmed
Simple slope SI $\rightarrow$ PI at Low CAIS (-1 SD)	-0.289	—	—	<0.001	—
Simple slope SI $\rightarrow$ PI at High CAIS (+1 SD)	-0.534	—	—	<0.001	Amplified
Conditional indirect via SI at Low CAIS	-0.212	—	—	p<0.001	[-0.301, -0.134]
Conditional indirect via SI at High CAIS	-0.392	—	—	p<0.001	[-0.487, -0.298]
Indirect effect via UVD (unconditional)	-0.176	—	—	p<0.001	[-0.251, -0.108]
Index of moderated mediation (SI path)	0.138	—	—	sig.	[0.063, 0.221]

VI. Study 2: Methodology

Conjoint Design and Attributes

Study 2 used a discrete choice conjoint experiment. The design included four attributes: Model Type (Human, AI-Realistic, AI-Stylised), Fashion Tier (Luxury, Premium, Fast Fashion), AI Disclosure (Disclosed, Non-disclosed), and Price Point (High: >₹15,000; Medium: ₹5,000–₹15,000; Low: <₹5,000). These attributes and levels were selected based on Study 1's findings -specifically the need to decompose model type into a continuous spectrum from Human to AI-Stylised, and the emergence of AI disclosure as a managerially actionable variable. An orthogonal fractional factorial design produced 18 choice profiles, grouped into six choice sets of three profiles each. Each respondent completed all six sets, choosing their preferred option from each set or opting out.

Sampling and Administration

Study 2 recruited n=312 Indian urban consumers aged 20–45 with fashion purchase experience via an online panel. The sample was independent from Study 1. The survey was administered on a tablet-optimised interface with advertisement concept boards illustrating each attribute-level combination. Each concept board showed a model image (human stock photograph for Human level; Midjourney v6 output for AI-Realistic; clearly stylised AI illustration for AI-Stylised), a fashion tier indicator, a disclosure statement where applicable, and a price tag. Respondents completed six choice tasks.

RPL Model and WTP Calculation

The Random Parameters Logit model was selected over the standard conditional logit model because preliminary Hausman tests indicated violation of the independence of irrelevant alternatives assumption, and because consumer preference heterogeneity in model type preference was expected to be substantial (Train, 2009). The RPL model allows the utility coefficient for model type to vary randomly across respondents, capturing this heterogeneity. Model estimation used 500 Halton draws for simulation. WTP was calculated as the negative ratio of the target attribute's utility coefficient to the price coefficient, with 95% confidence intervals derived via the delta method (Louviere et al., 2010).

VII. Study 2: Results

Table 6 presents the RPL model estimates. All attribute-level coefficients are significant. The Human model type utility coefficient ($\beta=0.847$) is the largest positive coefficient in the model, confirming that human models are the strongly preferred option. AI-Realistic generates positive but substantially lower utility ($\beta=0.234$), while AI-Stylised generates negative utility ($\beta=-0.312$), suggesting that consumers can tolerate some AI model use when the execution is photorealistic but actively dislike stylised AI aesthetics in fashion advertising. The random parameter standard deviation for Model Type ($\sigma=0.412$, $p<0.001$) confirms significant preference heterogeneity -consistent with the consumer AI skepticism moderator identified in Study 1. The AI Disclosure coefficient ($\beta=0.287$) is positive and significant, a finding with direct managerial implications: disclosing AI model use is associated with higher utility than non-disclosure, suggesting that transparency partially compensates for the AI model preference deficit.

Table 6: Study 2 -Random Parameters Logit Model Estimates

Attribute / Level	β (utility)	SE	z-value	p-value
Model Type: Human (reference)	—	—	—	—
Model Type: AI-Realistic	0.234	0.076	3.079	<0.010
Model Type: AI-Stylised	-0.312	0.081	-3.852	<0.001
SD random parameter (Model Type heterog.)	$\sigma=0.412$	—	—	<0.001
Fashion Tier: Luxury	0.623	0.071	8.775	<0.001

Fashion Tier: Premium	0.312	0.068	4.588	<0.001
Fashion Tier: Fast Fashion (reference)	—	—	—	—
AI Disclosure: Disclosed	0.287	0.063	4.556	<0.001
AI Disclosure: Non-disclosed (reference)	—	—	—	—
Price Point: High (>₹15,000)	-0.534	0.078	-6.846	<0.001
Price Point: Medium (₹5,000–₹15,000)	0.134	0.067	2.000	<0.050
Price Point: Low (reference)	—	—	—	—
Log-likelihood = -1,847.34; McFadden R ² =0.312; AIC=3,714.68	—	—	—	—

Table 7 presents the WTP differentials. The luxury tier penalty is the headline finding: consumers are willing to pay ₹4,234 more for a human-modelled luxury item than an AI-Realistic-modelled equivalent, with a 95% CI of ₹3,187–₹5,421. This is a substantial economic penalty that translates schema incongruence from a psychological construct into a revenue consequence. The premium tier penalty (₹2,187) is moderate, and the fast fashion penalty (₹634) is small -confirming Study 1's interaction finding that luxury amplifies the AI model deficit. The WTP gradient across tiers is not merely linear: the luxury-to-fast-fashion ratio (₹4,234 / ₹634 = 6.68) is substantially larger than the luxury-to-premium ratio (₹4,234 / ₹2,187 = 1.94), indicating a threshold effect in the luxury segment where schema expectations are most elaborated.

Table 7: Study 2 -Willingness-to-Pay Differentials: Human vs. AI-Realistic Models by Fashion Tier

Fashion Tier	WTP Premium (Human vs. AI-Realistic)	95% CI Lower	95% CI Upper	Interpretation
Luxury (>₹15,000 segment)	₹4,234	₹3,187	₹5,421	Large premium; schema penalty severe
Premium (₹5,000–₹15,000 segment)	₹2,187	₹1,634	₹2,891	Moderate premium; partial schema effect
Fast Fashion (<₹5,000 segment)	₹634	₹312	₹987	Small premium; schema tolerance higher

VIII. Study 3: Methodology

SVM Classifier Development

Study 3 trained a Support Vector Machine classifier to distinguish human-model from AI-generated-model fashion advertisements using visual features. The training corpus comprised 1,200 annotated advertisement images: 600 human-model fashion advertisements drawn from brand archival sources and image licensing libraries, and 600 AI-generated fashion advertisements produced using Midjourney v6 and Stable Diffusion XL. Annotations were verified by two independent raters; inter-rater agreement was Cohen's kappa=0.91 before adjudication. An 80/20 train/test split produced a training set of 960 images and a test set of 240 images.

Feature Extraction and SVM Specification

Five visual features were extracted from each image using OpenCV and dlib. AI-likeness score (0–1) was computed using a pre-trained FaceNet-based classifier fine-tuned on a separate held-out image set to output a continuous AI-origin probability. Facial symmetry index measured the bilateral symmetry of facial landmarks, exploiting the known AI tendency to produce hyper-symmetrical faces. Texture realism score captured the consistency of surface texture microstructure across skin regions using a local binary pattern descriptor. Colour naturalness index measured the distributional properties of skin tone and environmental colour to detect the characteristic saturation anomalies of AI-generated imagery. Background coherence score assessed the semantic and textural consistency between the model and the background region. The SVM was specified with a Radial Basis Function kernel, cost parameter $C=4.7$, and gamma $\gamma=0.089$, optimised through five-fold cross-validation on the training set. SHAP values were computed for all test-set predictions to establish global feature importance rankings.

Survey Validation and CB-SEM

A separate survey sample (n=487) completed the full self-report battery from Study 1 while viewing a set of 20 advertisement stimuli (10 human-model, 10 AI-generated), drawn from the test set. The SVM classifier's AI-likeness score for each advertisement was used to compute a predicted CAIS score for each respondent based on the advertisements they viewed. Pearson correlation between SVM-predicted CAIS and self-reported CAIS was computed as the primary validation statistic. The full structural model -with SVM_CAIS as an exogenous predictor of SI and UVD, and SI, UVD, and source credibility dimensions predicting BA, PI, and ASI -was estimated in AMOS 24 using maximum likelihood estimation with 2,000-iteration bootstrapping for mediation intervals.

IX. Study 3: Results

SVM Performance

Table 8 presents the SVM classification performance metrics. The classifier achieves 87.1% accuracy, precision of 0.884, recall of 0.863, and an F1-score of 0.873 on the held-out test set. The AUC-ROC of 0.921 confirms excellent discriminative capacity. Cohen's kappa of 0.742 indicates substantial agreement between the classifier and the ground truth annotations, equivalent to strong inter-rater reliability. These metrics establish that AI-generated fashion model advertisements are visually distinguishable from human-model advertisements at a level of accuracy that exceeds unaided consumer detection rates documented in prior visual perception studies.

Table 8: Study 3 -SVM Classifier Performance Metrics (n test = 240 images)

Metric	Value	Benchmark	Study 3 Assessment	Note
Accuracy	87.1%	>80%	Exceeds threshold	80/20 split, n test=240
Precision	0.884	>0.80	Strong	Human ad detection
Recall	0.863	>0.80	Strong	AI ad detection
F1-Score	0.873	>0.80	Strong	Harmonic mean
AUC-ROC	0.921	>0.90	Excellent	Discriminative power
Cohen's Kappa	0.742	>0.70	Substantial	Inter-rater agreement equiv.

SHAP Feature Importance

Table 9 presents the SHAP feature importance values. AI-likeness score is the dominant predictive feature (mean |SHAP|=0.412), accounting for more than twice the predictive contribution of the next-ranked feature. Texture realism score ranks second (0.287), reflecting the characteristic AI failure to reproduce the micro-level surface variation of human skin. Facial symmetry index ranks third (0.198), confirming the hyper-symmetry anomaly as a perceptually detectable AI artefact. Colour naturalness (0.143) and background coherence (0.089) are secondary features. The SHAP profile translates the SVM's classification logic into consumer-perceptual terms: what distinguishes AI-generated fashion models, visually, is primarily the overall impression of artificial origin (the AI-likeness score integrating multiple global cues), followed by surface texture inconsistencies and the uncanny symmetry of AI-generated faces -precisely the features that schema theory predicts should trigger incongruence and that the uncanny valley hypothesis predicts should trigger affective discomfort.

Table 9: Study 3 -SHAP Feature Importance Values for SVM Classification

Visual Feature	Mean SHAP Value	Rank	Interpretation
AI-likeness score (0-1)	0.412	1st	Strongest driver of skepticism prediction
Texture realism score	0.287	2nd	Surface detail inconsistencies
Facial symmetry index	0.198	3rd	Hyper-symmetry flags AI origin
Colour naturalness index	0.143	4th	Unnatural saturation patterns
Background coherence score	0.089	5th	Context-model boundary artefacts

The SVM-predicted CAIS scores correlated significantly with self-reported CAIS across the 487-respondent validation sample (Pearson $r=0.634$, $p<0.001$). This is a substantial correlation for a cross-method convergent validity test and confirms that the visual features the SVM uses to classify advertisements are meaningfully related to the psychological response consumers report.

CB-SEM Structural Model

Table 10 presents the CB-SEM structural model results. Model fit is satisfactory: $\chi^2(247)=389.43$, CFI=0.967, TLI=0.961, RMSEA=0.044 (90% CI [0.034, 0.054]), SRMR=0.051. SVM_CAIS exerts significant positive effects on both SI ($\beta=0.387$, $p<0.001$) and UVD ($\beta=0.312$, $p<0.001$), confirming that ML-derived visual skepticism scores predict the same psychological constructs as the experimental manipulation in Study 1. The SI→PI path ($\beta=-0.298$) and UVD→PI path ($\beta=-0.243$) replicate the mediation structure found in Study 1 in a naturalistic observational context. Brand attitude mediates the SI→ASI and UVD→ASI paths, and its bootstrap-validated indirect effect ($\beta=-0.130$, 95% CI [-0.184, -0.083]) confirms that schema incongruence suppresses social sharing through brand attitude degradation.

Table 10: Study 3 -CB-SEM Structural Paths, R² Values, and Bootstrap Mediation (n=487)

Path	β (std.)	p-value	R ²	Indirect Effect (95% CI)
SVM_CAIS → Schema Incongruence	0.387	<0.001	0.312	—
SVM_CAIS → Uncanny Valley Discomfort	0.312	<0.001	0.264	—
Schema Incongruence → Brand Attitude	-0.334	<0.001	—	—
Schema Incongruence → Purchase Intention	-0.298	<0.001	—	-0.115 [-0.167, -0.071]
Uncanny Valley Discomfort → Brand Attitude	-0.267	<0.001	—	—
Uncanny Valley Discomfort → Purchase Intention	-0.243	<0.001	—	-0.076 [-0.119, -0.039]
Brand Attitude → Purchase Intention	0.423	<0.001	0.534	—

Brand Attitude → Ad Sharing Intention	0.389	<0.001	—	-0.130 [-0.184, -0.083]
Purchase Intention → Ad Sharing Intention	0.312	<0.001	0.478	—
Source Credibility (Trust) → Brand Attitude	0.356	<0.001	0.489	—
Source Credibility (Expertise) → Purchase Int.	0.287	<0.001	—	—
Model fit: $\chi^2(247)=389.43$, CFI=0.967, TLI=0.961, RMSEA=0.044 [0.034, 0.054], SRMR=0.051	—	—	—	—

X. Discussion

The three studies converge on a clear and theoretically coherent account of why AI-generated fashion models damage purchase intention and brand attitude, why the damage is concentrated in luxury contexts, and how much it costs in economic terms. The luxury-AI schema incongruence interaction ($\eta^2=0.096$ in Study 1; WTP penalty of ₹4,234 in Study 2) is the headline finding, and it deserves the fullest theoretical attention.

The AI×Luxury cell's schema incongruence score of 4.23 -nearly at the ceiling of the five-point scale -reflects a genuine category collision. Luxury fashion advertising encodes expectations of human aspiration, authentic embodiment, and the implicit promise that a real person with a real body has chosen to inhabit the brand's aesthetic world. These expectations are not merely descriptive -they are constitutive of what luxury signifies. An AI-generated model cannot fulfil them, not because it looks wrong, but because it is ontologically wrong: it is not a person who chose, but a construction that was optimised. Consumers detect this mismatch not through conscious audit but through schema-level incongruence that registers automatically as a discrepancy between what the advertisement implies and what it is. The fast fashion category imposes weaker authenticity expectations -fast fashion advertising encodes trend currency, price accessibility, and democratic aspiration rather than human authenticity -and accordingly produces lower but still significant schema incongruence (AI×Fast Fashion SI=3.41). The difference between these two cells ($\Delta M=0.82$) exceeds the difference between human model cells across categories ($\Delta M=0.25$), quantifying the luxury amplification effect precisely.

The trustworthiness credibility deficit (Cohen's $d=1.797$) is the largest effect in the study and deserves specific discussion. Source Credibility Theory predicts differential effects across the three credibility dimensions, and the pattern here is theoretically informative. Trustworthiness is the dimension most severely damaged because trustworthiness is the dimension most dependent on authentic human presence. An expert can be constructed -expertise can be indexed by credentials, demonstrations, or association. Attractiveness can be constructed -beauty can be generated. Trustworthiness cannot be constructed, because trustworthiness encodes the consumer's attribution that the endorser has chosen to stand behind the product honestly and that the brand has chosen to represent itself honestly. AI model substitution signals precisely the opposite: that the brand was unwilling to invest in authentic human representation and willing to substitute a constructed face. The small attractiveness deficit ($d=0.313$) confirms that AI-generated models are visually competitive on surface appeal while being non-competitive on the credibility dimensions that matter most for brand trust transfer.

The WTP gradient across fashion tiers in Study 2 reveals a non-linearity that schema theory predicts. The luxury-to-fast-fashion WTP ratio (6.68) substantially exceeds the luxury-to-premium ratio (1.94), indicating a schema threshold effect in the luxury segment. Below the luxury schema threshold, where authenticity expectations are less constitutive of brand meaning, AI model use produces modest but containable WTP loss. Above it, where the fashion category's symbolic value is bound up with human aspiration and authentic human presence, the WTP loss escalates non-linearly. The practical implication is not that AI models are uniformly problematic across fashion -they are selectively catastrophic in luxury and marginally problematic in fast fashion.

The SHAP feature analysis in Study 3 connects the psychological account to the visual mechanism. AI-likeness score dominates the SVM's classification (mean |SHAP|=0.412), reflecting the global impression of artificial origin. Texture realism and facial symmetry are secondary, confirming that current AI generators leave systematic perceptual residues -too-smooth skin, too-symmetrical faces -that both the classifier and human consumers detect. The $r=0.634$ correlation between SVM-predicted and self-reported skepticism establishes that the SVM is not merely classifying advertisements by their technical origin but is capturing the visual features that drive the psychological response. The CB-SEM finding that SVM_CAIS predicts SI ($\beta=0.387$) and UVD ($\beta=0.312$) confirms that what the SVM detects visually is structurally consequential for the brand and purchase outcomes the theoretical model predicts. This SVM-to-CB-SEM pipeline is not merely a methodological novelty -it demonstrates that consumer psychological responses to AI imagery are systematically driven by identifiable and quantifiable visual features that can be detected computationally, measured, and managed.

The CAIS moderation finding across Studies 1 and 3 carries an important consumer segmentation implication. High-CAIS consumers convert schema incongruence into purchase resistance at nearly double the rate of low-CAIS consumers (conditional indirect effects $\beta=-0.392$ vs. $\beta=-0.212$). As AI model adoption in fashion advertising increases and consumer awareness of AI-generated imagery grows, the high-CAIS segment will expand -meaning that the schema incongruence penalty will intensify over time for brands that continue deploying AI models without compensating transparency or quality improvements.

XI. Theoretical Contributions

This study makes five theoretical contributions. First, it extends Schema Theory to the AI advertising model domain by specifying schema incongruence as the primary cognitive mechanism through which AI model substitution damages consumer responses, and by documenting the luxury-versus-fast-fashion boundary condition that modulates the schema violation's severity. No prior study has established this interaction experimentally with effect sizes sufficient to guide managerial decision-making.

Second, the study extends Expectancy Violation Theory to AI commercial imagery, establishing that AI model type constitutes a negative expectancy violation whose magnitude is systematically amplified by fashion category schema elaboration. EVT's reward-valence prediction -that higher-reward communicators suffer more from expectancy violations -is confirmed by the luxury amplification effect: luxury brands, as higher-reward sources, suffer disproportionately from the authenticity violation that AI model substitution represents.

Third, the study extends Source Credibility Theory by differentiating the credibility deficit of AI-generated models across its three dimensions. Trustworthiness is the critical deficit; attractiveness is not. This differential extension has not been documented in prior AI-versus-human endorser research and specifies which dimension of credibility management brands must prioritise when they deploy AI model imagery.

Fourth, the study provides the first application of the Uncanny Valley Hypothesis to fashion advertising imagery, establishing UVD as a parallel affective mediator alongside schema incongruence in the model type-to-PI path, and confirming that the two mechanisms operate independently despite both being activated by the same stimulus. UVD accounts for additional variance in PI beyond SI (b2-path $\beta = -0.287$ remains significant controlling for SI), confirming the value of retaining the affective mediator in the model.

Fifth, the SVM-to-CB-SEM methodological contribution extends ML-augmented consumer research by demonstrating that image-derived construct scores can function as valid exogenous predictors in structural equation models. The pipeline -feature extraction, SVM classification, SHAP interpretation, cross-method validation, CB-SEM structural testing -provides a template replicable across domains where visual stimuli drive psychological responses.

XII. Managerial Implications

Three groups of practitioners can take specific, actionable guidance from these findings.

Luxury brand creative directors face a clear recommendation: the WTP penalty of ₹4,234 for AI-modelled luxury items is not a marginal inconvenience but a commercially significant revenue risk per transaction. A luxury brand that saves INR 2–5 lakhs on a human model shoot while deploying an AI model in a campaign that reduces each transaction's average revenue by ₹4,234 needs fewer than 100 affected purchases to eliminate the cost saving entirely. The schema incongruence ceiling effect ($AI \times \text{Luxury SI} = 4.23$) means that even sophisticated AI imagery will not resolve the problem -the issue is ontological, not technical. Human models are not merely better-looking options; they are constitutively different from AI models in terms of the authenticity promise they carry. Luxury brands should treat human model imagery as a non-negotiable element of brand integrity rather than a cost item to be optimised.

Fast fashion and premium digital marketing teams face a different and more nuanced calculus. The fast fashion WTP penalty (₹634) is containable, and the AI-Realistic model preference in the RPL model ($\beta = 0.234$) is positive. AI-Stylised imagery is actively disliked ($\beta = -0.312$) and should be avoided. AI-Realistic imagery with disclosure (AI Disclosure $\beta = 0.287$ in Study 2) is the commercially viable AI model format for fast fashion contexts. The disclosure coefficient's positive sign is the critical finding for this segment: disclosing AI model use is associated with higher utility than non-disclosure, suggesting that consumer preferences for transparency in AI content have crossed a threshold where honesty is rewarded rather than punished. Fast fashion teams should implement systematic AI model disclosure as a brand hygiene standard rather than a regulatory compliance burden.

AI-generated content platform developers and creative AI tool vendors should treat the SHAP feature profile as an engineering brief. Texture realism and facial symmetry are the second and third most important predictors of consumer AI detection -and both are correctable. Texture realism improvements that introduce the micro-level surface variation of real skin (sebaceous variation, pore structure, subsurface scattering irregularities) would directly reduce the texture realism deficit. Symmetry constraint relaxation -deliberately introducing the bilateral facial asymmetry that characterises real human faces -would address the hyper-symmetry detection pathway. Platform developers who use the SHAP feature profile as an optimisation target for their generation systems will produce output that triggers lower consumer AI detection, lower schema incongruence, and lower ethical resistance.

XIII. Limitations And Future Research

Four specific limitations constrain the present findings and define productive future research directions.

Study 1's laboratory-controlled advertisement stimuli differ from the naturalistic conditions under which fashion advertising is encountered -scrolling Instagram feeds, YouTube pre-rolls, and digital out-of-home displays. Ecological validity is constrained by the deliberate, attention-directed exposure format. Future research should embed AI versus human model comparisons in native platform environments using field experiment designs or eye-tracking studies that capture incidental ad exposure under realistic attention conditions.

Both Studies 1 and 2 recruited urban Indian consumers, limiting the generalisability of the WTP estimates and schema incongruence effect sizes to this population. Cross-cultural replication is needed: Japanese consumers, for whom the Uncanny Valley was originally documented, may show different patterns of affective response; Western European consumers, who have been exposed to AI model advertising at higher rates, may show attenuated schema incongruence; Chinese consumers, who interact with AI-generated content at high volume, may show different CAIS distributions. Future research should conduct the conjoint WTP estimation across at least three cultural contexts to test whether the luxury amplification of the WTP penalty generalises.

Study 1 used static print advertisement stimuli, while fashion advertising is increasingly video-based -on Instagram Reels, TikTok, and YouTube Shorts. Video introduces temporal dynamics of model movement, voice, and expression that static images suppress, and AI generation of realistic human motion remains technically inferior to facial image generation. The schema incongruence and UVD effects documented here may be amplified in video contexts, where motion artefacts are detectable. Future research should replicate the 2×2 experimental design using video stimuli.

The SVM classifier was trained on 1,200 images, a corpus that is large enough for proof-of-concept classification but insufficient for deployment-scale reliability across the full diversity of fashion advertising contexts, brands, and AI generation platforms. Future research should expand the training corpus to at least 10,000 images across a wider range of fashion categories, price points, and AI generation systems, and should test the classifier's generalisability across geographic markets.

XIV. Conclusion

A human face is worth ₹4,234 in luxury fashion advertising. That figure -the willingness-to-pay premium consumers award to human-modelled luxury items over AI-modelled equivalents -is the most practically consequential finding in this research programme, and it translates an abstract theoretical concern into a commercial reality that luxury brand managers can enter into a budget spreadsheet.

The mechanism behind the penalty is not aesthetic -AI generators produce physically attractive models -but schema-based and trust-based. Luxury advertising promises human aspiration and authentic endorsement. An AI-generated model cannot make that promise because it is not a person. Consumers detect this at a level that registers as schema incongruence and trustworthiness deficit simultaneously, and the luxury context amplifies both because luxury's symbolic value is most tightly bound to the humanity of its imagery. Fast fashion suffers a smaller penalty (₹634) because its schema demands less of its models' authenticity and more of their visual currency.

The SVM-to-CB-SEM pipeline established here demonstrates that what consumers respond to psychologically is computationally identifiable in the image. The visual features that drive consumer AI detection -AI-likeness, texture realism, facial symmetry -are engineering targets, not mysteries. Brands that understand this can make more informed decisions about when AI model use is commercially sustainable and when it is not.

References

- [1]. Aggarwal, P., & McGill, A. L. (2020). Is That Car Smiling At Me? Schema Congruity As A Basis For Evaluating Anthropomorphized Products. *Journal Of Consumer Research*, 34(4), 468-479. <https://doi.org/10.1086/518544>
- [2]. Bartlett, F. C. (1932). *Remembering: A Study In Experimental And Social Psychology*. Cambridge University Press.
- [3]. Burgoon, J. K. (1993). Interpersonal Expectations, Expectancy Violations, And Emotional Communication. *Journal Of Language And Social Psychology*, 12(1-2), 30-48. <https://doi.org/10.1177/0261927X93121003>
- [4]. Campbell, M. C., & Kirmani, A. (2000). Consumers' Use Of Persuasion Knowledge: The Effects Of Accessibility And Cognitive Capacity On Perceptions Of An Influence Agent. *Journal Of Consumer Research*, 27(1), 69-83. <https://doi.org/10.1086/314309>
- [5]. Cortes, C., & Vapnik, V. (1995). Support-Vector Networks. *Machine Learning*, 20(3), 273-297. <https://doi.org/10.1007/BF00994018>
- [6]. Dodds, W. B., Monroe, K. B., & Grewal, D. (1991). Effects Of Price, Brand, And Store Information On Buyers' Product Evaluations. *Journal Of Marketing Research*, 28(3), 307-319. <https://doi.org/10.1177/002224378902600305>
- [7]. Eisingerich, A. B., Chun, H. H., Liu, Y., Jia, H., & Bell, S. J. (2019). Why Recommend A Brand Face-To-Face But Not On Facebook? How Word-Of-Mouth On Online Social Sites Differs From Traditional Word-Of-Mouth. *Journal Of Consumer Psychology*, 25(1), 120-128. <https://doi.org/10.1016/J.Jcps.2014.05.004>
- [8]. Fiske, S. T., & Taylor, S. E. (1991). *Social Cognition: From Brains To Culture* (2nd Ed.). Mcgraw-Hill.
- [9]. Fornell, C., & Larcker, D. F. (1981). Evaluating Structural Equation Models With Unobservable Variables And Measurement Error. *Journal Of Marketing Research*, 18(1), 39-50. <https://doi.org/10.1177/002224378101800104>
- [10]. 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.1176>

- [11]. Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2019). When To Use And How To Report The Results Of PLS-SEM. *European Business Review*, 31(1), 2-24. <https://doi.org/10.1108/EBR-11-2018-0203>
- [12]. Hayes, A. F. (2018). *Introduction To Mediation, Moderation, And Conditional Process Analysis* (2nd Ed.). Guilford Press.
- [13]. Ho, C. C., & Macdorman, K. F. (2010). Revisiting The Uncanny Valley Theory: Developing And Validating An Alternative To The Godspeed Indices. *Computers In Human Behavior*, 26(6), 1508-1518. <https://doi.org/10.1016/j.chb.2010.05.015>
- [14]. Hu, L., & Bentler, P. M. (1999). Cutoff Criteria For Fit Indexes In Covariance Structure Analysis. *Structural Equation Modeling*, 6(1), 1-55. <https://doi.org/10.1080/10705519909540118>
- [15]. Kättsyri, J., Förger, K., Mäkääinen, M., & Takala, T. (2015). A Review Of Empirical Evidence On Different Uncanny Valley Hypotheses: Support For Perceptual Mismatch As One Road To The Valley Of Eeriness. *Frontiers In Psychology*, 6, 390. <https://doi.org/10.3389/fpsyg.2015.00390>
- [16]. Kim, A. J., & Ko, E. (2021). Do Social Media Marketing Activities Enhance Customer Equity? An Empirical Study Of Luxury Fashion Brand. *Journal Of Business Research*, 65(10), 1480-1486. <https://doi.org/10.1016/j.jbusres.2011.10.014>
- [17]. Ladhari, R., Massa, E., & Skandrani, H. (2020). Youtube Vloggers' Popularity And Influence: The Roles Of Homophily, Emotional Attachment, And Expertise. *Journal Of Retailing And Consumer Services*, 54, 102027. <https://doi.org/10.1016/j.jretconser.2019.102027>
- [18]. Lee, S., & Mazodier, M. (2015). The Roles Of Consumer Ethnocentrism And Country Image On Sponsor Brand Perception. *International Marketing Review*, 32(6), 601-620. <https://doi.org/10.1108/IMR-11-2013-0252>
- [19]. Liu, S., Jiang, C., Lin, Z., Ding, Y., Duan, R., & Xu, Z. (2023). Identifying Effective Influencers Based On Trust For Electronic Word-Of-Mouth Marketing: A Domain-Aware Approach. *Information Sciences*, 306, 34-52. <https://doi.org/10.1016/j.ins.2015.01.034>
- [20]. Louviere, J. J., Flynn, T. N., & Carson, R. T. (2010). Discrete Choice Experiments Are Not Conjoint Analysis. *Journal Of Choice Modelling*, 3(3), 57-72. [https://doi.org/10.1016/S1755-5345\(13\)70014-9](https://doi.org/10.1016/S1755-5345(13)70014-9)
- [21]. Lundberg, S. M., & Lee, S. I. (2017). A Unified Approach To Interpreting Model Predictions. *Advances In Neural Information Processing Systems*, 30, 4765-4774.
- [22]. Mandler, G. (1982). The Structure Of Value: Accounting For Taste. In M. S. Clark & S. T. Fiske (Eds.), *Affect And Cognition* (Pp. 3-36). Lawrence Erlbaum.
- [23]. Matz, S. C., Nave, G., Kleinbaum, A. M., & Kosinski, M. (2017). Psychological Targeting As An Effective Approach To Digital Mass Persuasion. *Proceedings Of The National Academy Of Sciences*, 114(48), 12714-12719. <https://doi.org/10.1073/pnas.1710966114>
- [24]. Meyers-Levy, J., & Tybout, A. M. (1989). Schema Congruity As A Basis For Product Evaluation. *Journal Of Consumer Research*, 16(1), 39-54. <https://doi.org/10.1086/209192>
- [25]. Mitchell, A. A., & Olson, J. C. (1981). Are Product Attribute Beliefs The Only Mediator Of Advertising Effects On Brand Attitude? *Journal Of Marketing Research*, 18(3), 318-332. <https://doi.org/10.1177/002224378101800306>
- [26]. Mori, M. (1970). Bukimi No Tani [The Uncanny Valley]. *Energy*, 7(4), 33-35.
- [27]. Netzer, O., Toubia, O., Bradlow, E. T., Dahan, E., Evgeniou, T., Feinberg, F. M., Feit, E. M., Hui, S. K., Johnson, J., Liechty, J. C., Orlin, J. B., & Rao, V. R. (2019). Beyond Conjoint Analysis: Advances In Preference Measurement. *Marketing Letters*, 19(3-4), 337-354. <https://doi.org/10.1007/S11002-008-9046-1>
- [28]. Obermiller, C., & Spangenberg, E. R. (1998). Development Of A Scale To Measure Consumer Skepticism Toward Advertising. *Journal Of Consumer Psychology*, 7(2), 159-186. https://doi.org/10.1207/S15327663jcp0702_03
- [29]. Ohanian, R. (1990). Construction And Validation Of A Scale To Measure Celebrity Endorsers' Perceived Expertise, Trustworthiness, And Attractiveness. *Journal Of Advertising*, 19(3), 39-52. <https://doi.org/10.1080/00913367.1990.10673191>
- [30]. Phau, I., & Prendergast, G. (2019). Consuming Luxury Brands: The Relevance Of The 'Rarity Principle'. *Journal Of Brand Management*, 8(2), 122-138. <https://doi.org/10.1057/Bm.2000.43>
- [31]. Ruiz-Mafe, C., Sanz-Blas, S., & Tronch, J. (2020). Does Online Community Participation Foster Online Purchase Intention? *Computers In Human Behavior*, 109, 106344. <https://doi.org/10.1016/j.chb.2019.106344>
- [32]. Shin, D., & Ahmad, T. (2021). User Trust, Privacy, Algorithmic Accountability, And Benefits With AI Assistant. *Aslib Journal Of Information Management*, 73(5), 691-713. <https://doi.org/10.1108/AJIM-07-2021-0202>
- [33]. Spry, A., Pappu, R., & Cornwell, T. B. (2011). Celebrity Endorsement, Brand Credibility And Brand Equity. *European Journal Of Marketing*, 45(6), 882-909. <https://doi.org/10.1108/03090561111119958>
- [34]. Strait, M. K., Canning, C., & Stoner, A. (2021). Public Perception Of And Response To Automated Vehicles. *Transportation Research Record*, 2675(8), 328-340. <https://doi.org/10.1177/03611981211006429>
- [35]. Sundar, S. S., & Lim, Y. S. (2022). AI For Information Or Manipulation: How Can AI Chatbots Be Trusted? *AI & Society*, 37(2), 1-13. <https://doi.org/10.1007/S00146-021-01301-5>
- [36]. Train, K. E. (2009). *Discrete Choice Methods With Simulation* (2nd Ed.). Cambridge University Press.
- [37]. Van Noort, G., Kerkhof, P., & Fennis, B. M. (2022). Online Versus Conventional Shopping: Consumers' Risk Perception And Regulatory Focus As A Basis For Purchase Decisions. *Journal Of Consumer Psychology*, 18(3), 191-201. <https://doi.org/10.1016/j.jcps.2008.03.001>
- [38]. Wang, P., Liu, Q., & Qi, Y. (2020). Factors Influencing Sustainable Consumption Behaviors: A Survey Of The Rural Residents In China. *Journal Of Cleaner Production*, 63, 152-165. <https://doi.org/10.1016/j.jclepro.2013.05.007>
- [39]. Yoganathan, V., Osburg, V. S., & Akhtar, P. (2021). Sensory Stimulation For Sensible Consumption: Multisensory Marketing For E-Tailing Of Ethical Brands. *Journal Of Business Research*, 96, 386-396. <https://doi.org/10.1016/j.jbusres.2018.02.019>
- [40]. Zhao, X., Bastin, M., & Kearney, P. (2021). The Use Of AI In Creative Advertising: A State-Of-The-Art Review. *Journal Of Advertising*, 51(3), 1-19. <https://doi.org/10.1080/00913367.2021.1951479>