

The effect of error acknowledgment on consumer responses to human and AI influencers

İnsan ve Yapay Zekâ Influencer'larda hata üstlenme biçiminin tüketici tepkilerine etkisi

Selçuk Yasin Yıldız¹ 

Abstract

Human and artificial intelligence (AI) influencers increasingly share product recommendations, yet how consumers react when these actors respond to an informational error remains unclear. This study examines whether acknowledging responsibility for an error, rather than denying it, produces different consumer responses for human and AI influencers. A 2 (influencer type: human vs. AI) $\times$ 2 (error response: acknowledgment vs. denial) between-subjects scenario experiment was conducted with 496 social media users. Data were analysed using MANOVA, ANOVAs, and moderated serial mediation with 5,000 bootstrap samples. Error acknowledgment increased perceived warmth, parasocial interaction, trust, and purchase intention for both influencer types, although effects were substantially stronger for human influencers. A crossover interaction emerged for perceived competence: acknowledgment raised competence for the human influencer but lowered it for the AI influencer. Indirect effects on trust and purchase intention operated mainly through warmth, whereas the competence penalty reached behavioural intentions only serially via parasocial interaction. Findings specify boundary conditions of the pratfall effect and reveal a competence cost of AI transparency.

Keywords: AI Influencers, Error Acknowledgment, Parasocial Interaction

JEL Codes: M31

Öz

İnsan ve yapay zekâ (AI) influencer'lar ürün tavsiyelerini giderek daha yaygın biçimde paylaşmaktadır; ancak bu aktörlerin bir bilgi hatasına verdikleri yanıt karşısında tüketicilerin nasıl tepki verdiği belirsizliğini korumaktadır. Bu çalışma, hatanın sorumluluğunu üstlenmenin, hatayı reddetmeye kıyasla, insan ve yapay zekâ influencer'lar için farklı tüketici tepkileri üretip üretmediğini incelemektedir. 496 sosyal medya kullanıcısıyla 2 (influencer tipi: insan vs. AI) $\times$ 2 (hata yanıtı: üstlenme vs. reddetme) gruplar arası senaryo deneyi yürütülmüştür. Veriler MANOVA, ANOVA ve 5.000 bootstrap örneklemli düzenlenmiş seri aracılık analizleriyle çözümlenmiştir. Hata üstlenme her iki influencer tipinde de algılanan samimiyeti, parasosyal etkileşimi, güveni ve satın alma niyetini artırmış; ancak etkiler insan influencer'larda belirgin biçimde daha güçlü çıkmıştır. Algılanan yetkinlikte çaprazlama etkileşim gözlenmiştir: Üstlenme insan influencer'ın yetkinliğini artırırken yapay zekâ influencer'ın yetkinliğini düşürmüştür. Güven ve satın alma niyeti üzerindeki dolaylı etkiler ağırlıklı olarak samimiyet üzerinden işlerken, yetkinlik cezası davranışsal niyetlere yalnızca parasosyal etkileşim aracılığıyla seri biçimde ulaşmıştır. Bulgular, pratfall etkisinin sınır koşullarını belirlemekte ve yapay zekâ şeffaflığının yetkinlik maliyetini ortaya koymaktadır.

Anahtar Kelimeler: Yapay Zekâ Influencer'lar, Hata Üstlenme, Parasosyal Etkileşim

JEL Kodları: M31

¹ Assoc. Prof., Sivas Cumhuriyet University, Sivas, Türkiye, selcukyasinyl@cumhuriyet.edu.tr

ORCID: 0000-0002-1594-8799

Submitted: 1/08/2026

1st Revised: 16/09/2026

2nd Revised: 22/09/2026

Accepted: 24/09/2026

Online Published: 25/09/2026

Citation: Yıldız, S. Y., The effect of error acknowledgment on consumer responses to human and AI influencers, bmij (2026) 14 (3):1874-1891, doi: <https://doi.org/10.15295/bmij.v14i3.2844>

Introduction

Influencers, who occupy the center of engagement in social media marketing, have become powerful information sources that shape consumer decisions. A meta-analysis by Ao et al. (2023) reveals that source credibility exerts a stronger effect on purchase intention than any other influencer attribute examined. Credibility, in turn, operates largely through trust; the trust followers place in an influencer determines the persuasiveness of branded content (Lou & Yuan, 2019). In recent years, a fundamental innovation has been added to this picture: AI influencers (also referred to as virtual influencers) designed with generative artificial intelligence technologies are spreading rapidly, promising brands controllability, novelty, and scalability (Allal-Chérif et al., 2024; Sands et al., 2022; Thomas & Fowler, 2021).

Nevertheless, existing comparative evidence indicates that virtual influencers generally face perceptions of lower trustworthiness, authenticity, and parasocial connection than their human counterparts (Belanche et al., 2024; Franke et al., 2023; Zhang & Abdullah, 2026). Disclosing that content originates from artificial intelligence (AI disclosure) weakens brand authenticity and trust in most cases (Baek et al., 2026; Brüns & Meißner, 2024; Muniz et al., 2024), and in negative events, the way blame is attributed to virtual influencers differs from that of human influencers (Joel-Edgar et al., 2025). These findings suggest that AI transparency is not a one-sided virtue in the eyes of consumers and that its gains and costs must be evaluated together.

At the same time, influencers are not flawless information sources. Sharing inaccurate or incomplete information, undisclosed sponsorships, and similar transgressions are among the ordinary risks of the influencer ecosystem, and followers can punish such transgressions with credibility loss, unfollowing, and boycott threats (Cocker et al., 2021). The crisis communication literature shows that apology and corrective action can be effective, whereas silence or evasive responses can backfire (Eriksson, 2018). Almost all of this literature, however, is built on human actors. When an informational error emerges, whether the choice between openly acknowledging responsibility for the error and denying it by attributing it to external factors produces different outcomes depending on whether the actor is human or artificial intelligence has not been sufficiently examined experimentally.

This gap is also theoretically intriguing. The pratfall effect in social psychology suggests that small blunders by competent individuals can make them more likable; yet the empirical accumulation of this effect in marketing and consumer behavior contexts remains quite limited. The algorithm aversion literature, by contrast, offers a prediction in the opposite direction: people lose trust in algorithms that make a mistake far more quickly than in humans who make the same mistake (Berger et al., 2021; Dietvorst et al., 2015). Will an AI influencer's acknowledgment of its error therefore create a perception of human-like warmth, or will it violate the perfection standard expected of machines and lead to a loss of perceived competence? This study seeks to answer this question with a 2×2 between-subjects experimental design.

The purpose of the study is to examine the effects of influencer type (human vs. artificial intelligence) and error response (acknowledging responsibility vs. denying the error) on consumers' perceived warmth, competence, parasocial interaction, trust, and purchase intention, together with the mechanisms underlying these effects. The study aims to contribute to the literature by integrating the Stereotype Content Model, parasocial interaction theory, the algorithm aversion perspective, and the Elaboration Likelihood Model within a single model in which each framework accounts for a distinct link: the Stereotype Content Model identifies the evaluative dimensions (warmth and competence) through which consumers interpret an error response; the algorithm aversion perspective explains why these interpretations are expected to differ for AI influencers; parasocial interaction theory explains how these evaluations are carried to relational and behavioral outcomes; and the Elaboration Likelihood Model, operationalized through need for cognition, specifies for whom error acknowledgment should matter most. For terminological consistency, the term AI influencer is used throughout for the computer-generated influencer examined in this study, whereas the label virtual influencer is retained only when prior studies are reported in their own terms; likewise, error acknowledgment refers to the influencer's explicit acceptance of responsibility for the error, and error denial refers to its rejection of responsibility through attribution to external factors. The following section develops the theoretical framework and hypotheses; the methodology, findings, and conclusion and discussion sections are then presented.

Literature review and hypothesis development

The stereotype content model: Warmth and competence

The Stereotype Content Model (SCM) proposes that social perception is organized along two fundamental dimensions: warmth, reflecting intentions, and competence, reflecting the capacity to enact those intentions (Cuddy et al., 2008; Fiske, 2018; Fiske et al., 2007). There is an asymmetric priority between the two dimensions; because individuals first assess whether the other party is friend or foe, warmth information is processed earlier in impression formation and predicts global evaluations more strongly (Abele et al., 2021; Brambilla et al., 2012). The model has also been successfully transferred to marketing; brand and firm perceptions have been shown to form along the same two dimensions (Aaker et al., 2010; Gidaković & Žabkar, 2022; Kolbl et al., 2019), and both dimensions have been shown to nurture trust toward artificial actors ranging from chatbots to voice assistants (Cheng et al., 2022; Sung et al., 2023). Current evidence in the virtual influencer context likewise indicates that consumers stereotype human-like virtual influencers along the same dimensions and that warmth comes to the fore in willingness to follow recommendations (El Hedhli et al., 2023). Applied to the present context, the SCM implies that consumers treat an influencer's response to an error as diagnostic information on both dimensions: the way the influencer handles the error conveys its intentions toward followers (warmth), whereas what the response reveals about the origin of the error conveys its ability to provide accurate recommendations (competence). Accordingly, the present study positions warmth and competence as the two evaluative mechanisms through which error acknowledgment is expected to shape consumer responses, and the hypotheses below specify how influencer type alters each of these inferences.

Error acknowledgment, apology, and perceived warmth

Organizational and consumer research demonstrates that taking responsibility for an error is a powerful social signal. Responsibility-focused apologies convey remorse, empathy, and goodwill, which increases perceived warmth and supports trust repair (Chung & Lee, 2021; Yang et al., 2025). In experimental studies of service failure, apologizing actors have been found warmer and more likable (Mahmood & Huang, 2024). There are, however, important indications that this effect depends on the actor's human qualities: in robot service failures, apology increased warmth perceptions only for humanoid robots (Choi et al., 2021), and human employees' apologies were perceived as more sincere than robot apologies, with this sincerity gap determining recovery satisfaction (Hu et al., 2021). This evidence links the SCM to the first hypothesis in two steps. First, accepting responsibility is an other-oriented behavior that signals good intentions toward followers; error acknowledgment should therefore increase perceived warmth relative to error denial. Second, a warmth inference of this kind requires consumers to attribute genuine intentions to the source, and such attributions are made less readily for artificial than for human agents (Hu et al., 2021; Joel-Edgar et al., 2025); consequently, the warmth gain produced by acknowledgment should be smaller for an AI influencer than for a human influencer:

H1: *The positive effect of error acknowledgment (vs. error denial) on perceived warmth is stronger for a human influencer than for an AI influencer.*

Algorithm aversion and the competence penalty

The core finding of the algorithm aversion literature is that people lose trust in algorithms that err far more quickly than in humans who make the same mistake (Dietvorst et al., 2015). Two mechanisms stand out behind this asymmetry. The first is the perfection standard applied to machines: users expect automation to be flawless while regarding human error as natural, and a visible AI error causes a sharp drop in competence perceptions because it violates high accuracy expectations (Han & Ko, 2025; Molina & Sundar, 2024). The second is the belief that humans can learn from their mistakes whereas algorithms cannot improve; this perception renders an error more diagnostic of future performance (Berger et al., 2021). Visible AI errors consistently lower competence judgments (Toader et al., 2020; Zerilli et al., 2022), and aversion intensifies particularly in subjective tasks requiring human skill (Castelo et al., 2019). Conversely, some evidence shows that consumers can respond more leniently to brand crises when the source of the error is an algorithm (Srinivasan & Sarial-Abi, 2021), suggesting that artificial actors may be relatively protected in conditions where the error is not openly owned. In the present study, algorithm aversion is therefore used as the theoretical rationale for predicting how the competence dimension of the SCM responds to an AI influencer's error response, rather than as a construct to be measured. When an AI influencer explicitly acknowledges the error, it confirms that the error originated in its own functioning; given the perfection standard applied to machines and the presumed inability of algorithms to learn, this confirmation should be read as diagnostic of low ability and should reduce

perceived competence. When the AI influencer instead attributes the error to external factors, the error is less clearly tied to its own capabilities, so competence judgments should be comparatively preserved. The hypothesis is accordingly formulated in terms of perceived competence, the variable measured in the study:

H2: *For an AI influencer, error acknowledgment (vs. error denial) decreases perceived competence.*

Parasocial interaction

Parasocial interaction (PSI) refers to the one-sided yet friendship-like bond of intimacy that audiences form with media personalities and is one of the fundamental persuasion mechanisms of influencer marketing; strong parasocial bonds consistently increase purchase intention (Folkvord et al., 2020; Masuda et al., 2022; Sokolova & Kefi, 2020). The CASA paradigm, which proposes that computers are perceived as social actors, shows that people apply human social scripts to machines even with minimal social cues (Lombard & Xu, 2021; Nielsen et al., 2022); parasocial bonds with virtual influencers are therefore possible. Indeed, although direct comparisons do not always find differences in overall PSI levels between virtual and human influencers, they reveal that the mechanisms through which the bond is formed differ: low perceived humanness and similarity suppress PSI for virtual influencers (Stein et al., 2024), a lack of authenticity and perceptions of uncanniness can weaken the parasocial bond and persuasive power (Lou et al., 2023), while anthropomorphism can strengthen this bond (Dabiran et al., 2024; Mrad et al., 2025). Linking this theory to the present design, error acknowledgment is a relational cue that signals openness and accountability toward followers and should therefore strengthen the parasocial bond relative to error denial. Because PSI is nourished largely by relational cues such as warmth, authenticity, and credibility (Ashraf et al., 2023; Yuan & Lou, 2020), and because such cues are discounted when the source is perceived as less human-like (Lou et al., 2023; Stein et al., 2024), the parasocial gain from acknowledgment is expected to be smaller for an AI influencer:

H3: *The positive effect of error acknowledgment (vs. error denial) on parasocial interaction is stronger for a human influencer than for an AI influencer.*

Mediating mechanisms: Warmth, competence, and parasocial interaction

The SCM literature demonstrates that warmth and competence are not merely descriptive perceptions; they operate as mediating mechanisms through which brand and actor behaviors are transmitted to trust and purchase intention. The effects of diverse stimuli such as brand disclosure, advertising receptivity, and chatbot communication style on purchase intention are carried through warmth and competence perceptions (Ge et al., 2025; Roy & Naidoo, 2021; Shi & Jiang, 2023); even distal cues such as country stereotypes can influence intentions only when transferred to brand warmth and competence (Diamantopoulos et al., 2021). At the same time, warmth and competence are also antecedents of the parasocial bond; in the virtual influencer context, anthropomorphic warmth cues nurture willingness to follow recommendations primarily through warmth perceptions (El Hedhli et al., 2023). Within this framework, error acknowledgment is expected to influence trust and purchase intention through two complementary routes. In the first route, warmth and competence operate as parallel mediators that carry the effect of error acknowledgment to each outcome, because judgments of an influencer's intentions and ability are the bases on which followers decide whether to rely on its recommendations. Because this route involves two mediators and two outcomes, it is specified in separate sub-hypotheses:

H4a: *Perceived warmth mediates the effect of error acknowledgment (vs. error denial) on consumer trust.*

H4b: *Perceived competence mediates the effect of error acknowledgment (vs. error denial) on consumer trust.*

H4c: *Perceived warmth mediates the effect of error acknowledgment (vs. error denial) on purchase intention.*

H4d: *Perceived competence mediates the effect of error acknowledgment (vs. error denial) on purchase intention.*

In the second route, the warmth and competence inferences formed after the error response shape the parasocial bond, which in turn drives purchase intention; parasocial interaction is thus expected to act as a second-stage mediator that translates person perceptions into behavioral intentions (El Hedhli et al., 2023; Sokolova & Kefi, 2020). Because influencer type is expected to alter the first-stage effects of error acknowledgment (H1–H3), these serial indirect effects are estimated separately for the human and AI influencers:

H5a: Error acknowledgment (vs. error denial) influences purchase intention serially through perceived warmth and parasocial interaction (error acknowledgment → perceived warmth → parasocial interaction → purchase intention).

H5b: Error acknowledgment (vs. error denial) influences purchase intention serially through perceived competence and parasocial interaction (error acknowledgment → perceived competence → parasocial interaction → purchase intention).

The moderating role of need for cognition

The Elaboration Likelihood Model (ELM) proposes that persuasion proceeds via two routes: a central route (detailed processing of message arguments) and a peripheral route (heuristic processing based on simple cues) (Chaiken, 1980; Petty & Cacioppo, 1986). The likelihood that an individual follows either route depends partly on a stable disposition, need for cognition (NFC), defined as the tendency to engage in and enjoy effortful thinking (Cacioppo & Petty, 1982). Because the processing route actually used cannot be observed directly in a single-exposure scenario experiment, the present study uses NFC as an individual-difference indicator of elaboration propensity rather than as a direct measure of central versus peripheral processing. Individuals high in NFC are sensitive to argument quality, whereas those low in NFC respond more strongly to peripheral cues of a source or signal nature (Cacioppo et al., 1983); for example, promotion signals that involve no real price cut can mobilize only consumers low in NFC (Inman et al., 1990). An influencer's acknowledgment of an error can likewise function as a source-related cue that is independent of the quality of the recommendation itself. Consumers low in NFC are therefore expected to rely more heavily on this cue when forming trust and purchase intentions, whereas consumers high in NFC are expected to give greater weight to the substance of the recommendation and less weight to the acknowledgment cue. Because two outcome variables are involved, the moderation hypothesis is specified separately for each:

H6a: Need for cognition moderates the effect of error acknowledgment (vs. error denial) on consumer trust, such that the effect is weaker for consumers high (vs. low) in need for cognition.

H6b: Need for cognition moderates the effect of error acknowledgment (vs. error denial) on purchase intention, such that the effect is weaker for consumers high (vs. low) in need for cognition.

In line with the developed hypotheses, the conceptual model of the research (Figure 1) specifies error response (acknowledgment vs. denial) as the focal experimental factor and influencer type (human vs. AI) as the experimental factor that conditions its effects on perceived warmth, perceived competence, and parasocial interaction (H1–H3). Perceived warmth and perceived competence serve as parallel mediators (H4a–H4d), parasocial interaction serves as the second-stage mediator in the serial paths to purchase intention (H5a–H5b), consumer trust and purchase intention are the outcome variables, and need for cognition moderates the effect of error acknowledgment on both outcomes (H6a–H6b).

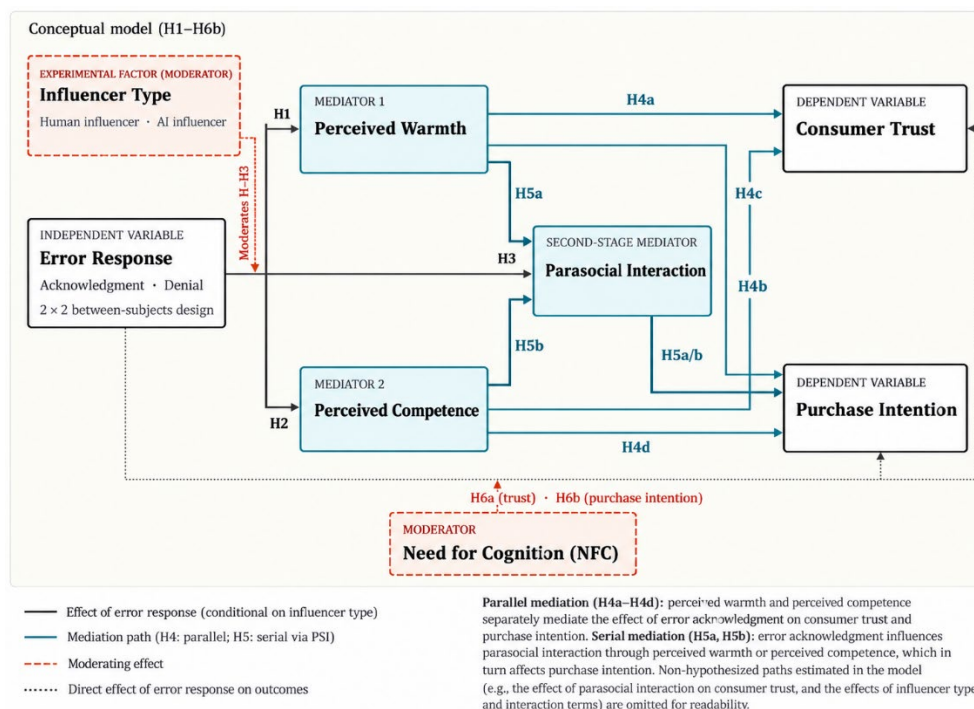


Figure 1: Conceptual model

Methodology

Research design

To test the effects of influencer type (human vs. artificial intelligence) and error response (acknowledging responsibility vs. denying the error/attributing it to external factors) on consumer perceptions and intentions, a 2 (Influencer Type: Human vs. AI) $\times$ 2 (Error Response: Acknowledgment vs. Denial) between-subjects scenario-based (vignette) experimental design was employed. Participants were randomly assigned to only one of the four experimental conditions. The sample size was determined through an a priori power analysis conducted with G*Power 3.1 (Test family: F tests; Statistical test: ANOVA: Fixed effects, special, main effects and interactions; numerator df = 1). With a medium effect size ($f = 0.25$), $\alpha = 0.05$, and power = 0.80, the calculated minimum sample size was 128. However, to account for potential data loss during the screening process and to ensure robust statistical power across experimental cells, a higher target of 80–100 participants per cell (320–400 in total) was established.

Participants and procedure

Data were collected from individuals over 18 years of age residing in Türkiye, who actively use at least one social media platform, via the Qualtrics online survey platform between 08.03.2025 and 26.07.2025. To reach a broad audience, the study adopted a non-hierarchical, non-probability sampling approach combining convenience sampling, which facilitates the participation of prospective respondents, and snowball sampling, which leverages network effects. Participation was entirely voluntary, and no monetary/financial incentive was offered in order to protect data quality and sincerity.

The data collection process yielded a total of 514 participants. In the data cleaning stage, a speeding screen was first applied based on survey completion times (median = 277 seconds); no participant completed the survey in less than one-third of the median duration. Next, 18 participants who gave the same response to all 19 scale items (straight-liners) were removed from the dataset, and no multivariate outliers were detected according to Mahalanobis distance ($p < .001$). The final analytical sample consists of 496 participants (Human–Acknowledgment: $n = 125$; AI–Acknowledgment: $n = 121$; Human–Denial: $n = 123$; AI–Denial: $n = 127$). Of the participants, 52.4% were female ($n = 260$) and 44.4% were male ($n = 220$), while 3.2% ($n = 16$) preferred not to state their gender; the mean age was 29.08 ($SD = 7.32$; range: 18–60). Of the participants, 62.3% held a bachelor's degree and 14.5% a postgraduate degree, and 61.5% stated that they had previously encountered AI influencers on social media but did not follow them.

After the informed consent and screening questions, participants examined the social media profile image and the influencer description belonging to the condition to which they were assigned. In the human condition, the profile was introduced as "a human social media influencer in their 30s producing content in the fields of travel, technology, and lifestyle"; in the AI condition, it was introduced as "a virtual artificial intelligence influencer designed with generative AI algorithms." In the experimental stimulus (vignette), a scenario was constructed in which the influencer recommended a popular mobile application/service in the travel and lifestyle category to their followers but made a minor informational error regarding the application's usage features. In the error response manipulation, the influencer either explicitly acknowledged responsibility for this informational error (error acknowledgment) or denied the error by attributing it to external factors (error denial). Ethics committee approval for this study was received from the Sivas Cumhuriyet University Social and Human Sciences Ethics Committee on 25 November 2024 (document no. 494383).

Measures

All constructs were measured with 7-point Likert-type scales (1 = Strongly disagree, 7 = Strongly agree). Perceived warmth (3 items) and perceived competence (3 items) were adapted from Bernritter et al. (2016) and Cuddy et al. (2008) on the basis of the Stereotype Content Model; parasocial interaction (PSI; 4 items) from Auter (1992) and Yılmazdoğan et al. (2021); and consumer trust (3 items) and purchase intention (3 items) from Dodds et al. (1991) and Zeithaml (1988). Need for cognition (NFC; 3 items), adapted from Cacioppo et al. (1984) and Dedeoğlu et al. (2021), was measured as the moderator variable and used as an individual-difference indicator of elaboration propensity within the ELM framework.

To test the effectiveness of the manipulations, participants were asked two manipulation check items concerning whether the profile examined was an AI influencer and whether the content creator acknowledged responsibility for the error, along with one realism check item concerning the credibility of the scenario.

Data analysis

Analyses were conducted in the R environment with the lavaan, semTools, psych, and emmeans packages. The first stage comprised data quality screens and manipulation checks; the second stage comprised confirmatory factor analysis (CFA; MLR robust estimator), reliability, convergent and discriminant validity, and measurement invariance tests across the four experimental groups; the third stage comprised two-way MANOVA and univariate 2×2 ANOVAs with simple effects analyses via emmeans; and the final stage tested conditional indirect effects and indices of moderated mediation through a path model built on observed composite scores (5,000 bootstrap samples, 95% percentile confidence intervals). Harman's single-factor test was applied for common method bias; the variance explained by a single factor was 39.7%, below the 50% threshold (Podsakoff et al., 2003). To test the moderating role of Need for Cognition (NFC) (H6a and H6b), regression models were estimated. The models included the error response condition (Acknowledgment vs. Denial, dummy coded), mean-centered NFC, and their interaction term as predictors. Ordinary Least Squares (OLS) standard errors were utilized. The degrees of freedom for the models were $F(3, 492)$. For the simple slopes analyses, conditional effects were estimated at ± 1 standard deviation of the mean-centered NFC ($SD = 1.201$)

Findings

Manipulation and realism checks

Two-way ANOVA results showed that both manipulations were successful. Participants in the AI condition perceived the profile as artificial intelligence at a significantly higher level ($M_{AI} = 6.53$ vs. $M_{Human} = 1.48$; $F(1, 492) = 6997.11$; $p < .001$; $\eta_p^2 = .934$), and participants in the acknowledgment condition reported that the error was acknowledged at a significantly higher level ($M_{Acknowledgment} = 6.52$ vs. $M_{Denial} = 1.50$; $F(1, 492) = 6382.04$; $p < .001$; $\eta_p^2 = .928$). Cross-over effects and interactions were non-significant (all $p > .17$). Scenario realism was high and equivalent across all four conditions ($M = 5.18$ – 5.34 ; all $p > .19$), indicating that the credibility of the scenarios did not differ across conditions.

Measurement model

The six-factor measurement model showed excellent fit to the data: $\chi^2(137) = 197.236$; $p = .001$; $\chi^2/df = 1.44$; CFI = .989; TLI = .986; RMSEA = .030 (90% CI: .020–.038); SRMR = .039 (Hu & Bentler, 1999). All standardized factor loadings ranged between .634 and .888 and were significant. As shown in Table 1, Cronbach's alpha and composite reliability (CR) values exceeded the .80 threshold and AVE values exceeded the .50 threshold for all constructs; convergent validity was established (Fornell & Larcker, 1981).

Table 1: Measurement Model Results

Construct	Items	Factor loadings	α	CR	AVE
Perceived warmth (WARM)	3	.833–.888	.900	.900	.751
Perceived competence (COMP)	3	.748–.846	.834	.834	.627
Parasocial interaction (PSI)	4	.634–.814	.841	.850	.580
Consumer trust (TRST)	3	.783–.882	.867	.870	.690
Purchase intention (PI)	3	.804–.867	.877	.878	.705
Need for cognition (NFC)	3	.712–.786	.804	.805	.579

Note: α = Cronbach's alpha; CR = composite reliability; AVE = average variance extracted. CFA was conducted with the MLR estimator ($N = 496$).

Discriminant validity was examined with both the Fornell–Larcker criterion and the HTMT ratio (Table 2). The square root of each construct's AVE is greater than its correlations with the other constructs, with a minor exception for the PSI–WARM pair (.762 vs. .770). However, all HTMT values are well below the .85 threshold (the highest value being .807 for TRST–PI; Henseler et al., 2015), establishing discriminant validity primarily through the HTMT criterion for this specific pair. Discriminant validity was established.

Table 2: Discriminant Validity: Fornell–Larcker Criterion and HTMT Ratios

	WARM	COMP	PSI	TRST	PI	NFC
WARM	(.867)	.490	.776	.610	.558	.118
COMP	.494	(.792)	.458	.369	.326	.158
PSI	.770	.480	(.762)	.589	.644	.091
TRST	.610	.372	.551	(.830)	.807	.328
PI	.560	.330	.611	.800	(.839)	.280
NFC	.119	.160	.085	.324	.280	(.761)

Note: Diagonal values (in parentheses) show the square root of the AVE; values below the diagonal show latent variable correlations; values above the diagonal show HTMT ratios.

Measurement invariance

To ensure the meaningfulness of between-group comparisons, measurement invariance was tested across the four experimental cells (Table 3). Metric invariance was fully supported ($\Delta\chi^2(39) = 30.641$; $p = .828$; $\Delta CFI = +.003$). The full scalar model showed deterioration ($\Delta CFI = -.013$; $p < .001$); when the intercepts of items PSI_4 and COMP_1 were freed in line with the modification indices, partial scalar invariance was established ($\Delta\chi^2(33) = 36.488$; $p = .310$; $\Delta CFI = -.001$; $|\Delta CFI| \leq .010$ criterion; Cheung & Rensvold, 2002). These findings indicate that between-group mean comparisons are defensible.

Table 3: Measurement Invariance Tests (Four Experimental Groups)

Model	CFI	RMSEA	ΔCFI	$\Delta\chi^2$ (df)	p
Configural	.960	.044	—	—	—
Metric	.963	.041	+.003	30.641 (39)	.828
Full scalar	.951	.046	-.013	80.496 (39)	< .001
Partial scalar	.962	.040	-.001	36.488 (33)	.310

Note: CFI and RMSEA are robust values. Δ values are relative to the preceding model (configural for metric; metric for the scalar models). In the partial scalar model, the intercepts of PSI_4 and COMP_1 were freed.

Main and interaction effects

Cell means and standard deviations are presented in Table 4. Levene's tests indicated that homogeneity of variance was met except for WARM ($F = 2.873$; $p = .036$); because Box's M test was significant ($\chi^2(45) = 79.95$; $p = .001$), Pillai's Trace, which is robust to this violation, was reported for the MANOVA results. The balanced cell sizes limit the impact of these violations.

Table 4: Means and Standard Deviations by Experimental Condition

Variable	Human–Acknowledgment	Human–Denial	AI–Acknowledgment	AI–Denial
Perceived warmth	5.83 (1.18)	3.44 (1.00)	4.02 (1.27)	3.14 (1.19)
Perceived competence	5.20 (1.31)	4.02 (1.08)	3.52 (1.05)	4.47 (1.21)
Parasocial interaction	5.21 (1.10)	3.83 (1.06)	3.69 (0.96)	3.06 (1.04)
Consumer trust	5.63 (0.92)	3.67 (1.20)	3.75 (1.20)	3.24 (1.15)
Purchase intention	5.55 (0.94)	3.48 (1.19)	3.50 (1.05)	2.93 (1.14)

Note: Values are mean (standard deviation). $n = 125$ (Human–Acknowledgment), 123 (Human–Denial), 121 (AI–Acknowledgment), 127 (AI–Denial).

The two-way MANOVA showed that influencer type (Pillai = .389; $F(5, 488) = 62.04$; $p < .001$), error response (Pillai = .474; $F(5, 488) = 87.91$; $p < .001$), and their interaction (Pillai = .262; $F(5, 488) = 34.67$; $p < .001$) were significant at the multivariate level. Univariate ANOVA results are presented in Table 5.

Table 5: Univariate 2 × 2 ANOVA Results

Dependent variable	Influencer type F (ηp^2)	Error response F (ηp^2)	Interaction F (ηp^2)
Perceived warmth	102.09*** (.172)	246.95*** (.334)	51.93*** (.095)
Perceived competence	34.67*** (.066)	1.28 (.003)	103.87*** (.174)
Parasocial interaction	148.08*** (.231)	115.04*** (.190)	16.03*** (.032)
Consumer trust	131.63*** (.211)	149.82*** (.233)	51.42*** (.095)
Purchase intention	176.89*** (.264)	182.84*** (.271)	59.50*** (.108)

Note: *** $p < .001$. Degrees of freedom are (1, 492) for all tests. Type III sums of squares were used.

Simple effects analyses (Holm-adjusted) revealed the nature of the interactions. Error acknowledgment very strongly increased warmth ($d = 2.19$), parasocial interaction ($d = 1.27$), trust ($d = 1.84$), and purchase intention ($d = 1.93$) in the human influencer condition; the same effects were significant but markedly weaker in the AI condition ($d = 0.72, 0.63, 0.43$, and 0.52 , respectively; all $p < .001$). This pattern supports H1 and H3.

The most striking finding is the full crossover interaction observed for perceived competence: while acknowledging the error increased the human influencer's perceived competence ($M_{\text{diff}} = 1.19$; $d = 0.99$; $p < .001$), it significantly decreased the AI influencer's perceived competence ($M_{\text{diff}} = -0.95$; $d = -0.84$; $p < .001$). In other words, error acknowledgment functioned as a positive competence cue for the human influencer but as a competence penalty for the AI influencer. The decrease observed in the AI condition supports H2; this pattern is consistent with the algorithm aversion account, although algorithm aversion itself was not measured directly.

Mediation and moderated mediation analyses

Coding the experimental conditions with dummy variables (Acknowledgment = 1; AI = 1) and an interaction term, a path model in which warmth and competence served as parallel mediators and parasocial interaction as a second-stage mediator was tested with 5,000 bootstrap samples. The model accounted for variance in warmth ($R^2 = .452$), competence ($R^2 = .220$), parasocial interaction ($R^2 = .513$), trust ($R^2 = .439$), and purchase intention ($R^2 = .498$). Standardized direct paths are presented in Table 6.

Table 6: Standardized Path Coefficients for the Structural Model

Path	β	p	Path	β	p
WARM $\leftarrow$ Acknowledgment	.765	< .001	TRST $\leftarrow$ WARM	.127	.027
WARM $\leftarrow$ AI	-.097	.030	TRST $\leftarrow$ COMP	.060	.212
WARM $\leftarrow$ Ack $\times$ AI	-.413	< .001	TRST $\leftarrow$ PSI	.119	.015
COMP $\leftarrow$ Acknowledgment	.451	< .001	TRST $\leftarrow$ Acknowledgment	.488	< .001
COMP $\leftarrow$ AI	.171	.002	TRST $\leftarrow$ AI	-.112	.040
COMP $\leftarrow$ Ack $\times$ AI	-.697	< .001	TRST $\leftarrow$ Ack $\times$ AI	-.305	< .001
PSI $\leftarrow$ WARM	.472	< .001	PI $\leftarrow$ WARM	-.050	.358
PSI $\leftarrow$ COMP	.141	< .001	PI $\leftarrow$ COMP	.014	.745
PSI $\leftarrow$ Acknowledgment	.113	.031	PI $\leftarrow$ PSI	.248	< .001
PSI $\leftarrow$ AI	-.276	< .001	PI $\leftarrow$ Acknowledgment	.601	< .001
PSI $\leftarrow$ Ack $\times$ AI	.042	.475	PI $\leftarrow$ AI	-.118	.021
			PI $\leftarrow$ Ack $\times$ AI	-.387	< .001

Note: WARM = perceived warmth; COMP = perceived competence; PSI = parasocial interaction; TRST = consumer trust; PI = purchase intention. Acknowledgment (1 = error acknowledgment, 0 = denial) and AI (1 = artificial intelligence, 0 = human) are dummy variables.

In the path model, the Acknowledgment $\times$ AI interaction for parasocial interaction was not significant after warmth and competence were included ($\beta = .042$; $p = .475$).

Conditional indirect effects and indices of moderated mediation are presented in Table 7. The indirect effect of error acknowledgment on trust through warmth was significant in both the human ($b = .282$; 95% CI [.032; .544]) and the AI condition ($b = .105$; 95% CI [.012; .215]). The index of moderated mediation was negative (IMM = $-.177$; 95% CI $[-.361; -.019]$), indicating a smaller indirect effect in the AI condition. The corresponding indirect effects through competence were not significant in either condition. Thus, H4a was supported and H4b was not. For purchase intention, the simple indirect effects through warmth were not significant for either human ($b = -.112$; 95% CI $[-.345; .127]$) or AI influencers ($b = -.042$; 95% CI $[-.136; .048]$). The simple indirect effects through competence were also not significant for either human ($b = .019$; 95% CI $[-.093; .139]$) or AI influencers ($b = -.015$; 95% CI $[-.113; .074]$). Accordingly, H4c and H4d were not supported. The serial indirect paths through parasocial interaction are reported separately under H5a and H5b.

The serial indirect effect of error acknowledgment on purchase intention through warmth and parasocial interaction was significant in both conditions (Human: $b = .264$; 95% CI [.165; .383]; AI: $b = .098$; 95% CI [.052; .158]). Its index of moderated mediation was negative (IMM = $-.166$; 95% CI $[-.258; -.095]$), supporting H5a. For the competence → parasocial interaction serial path, the indirect effect was positive for the human influencer ($b = .046$; 95% CI [.018; .083]) and negative for the AI influencer ($b = -.037$; 95% CI $[-.069; -.014]$). The AI-minus-human bootstrap contrast was also negative (IMM = $-.084$; 95% CI $[-.150; -.033]$), supporting H5b and confirming a difference between conditions.

Table 7: Conditional Indirect Effects and Indices of Moderated Mediation (5,000 Bootstrap Samples)

Indirect path	Condition	b	95% CI	Result
Ack → WARM → TRST	Human	.282	[.032; .544]	Significant
	AI	.105	[.012; .215]	Significant
Ack → COMP → TRST	Human	.078	$[-.044; .208]$	Non-significant
	AI	-.063	$[-.174; .034]$	Non-significant
Ack → WARM → PI	Human	-.112	$[-.345; .127]$	Non-significant
	AI	-.042	$[-.136; .048]$	Non-significant
Ack → COMP → PI	Human	.019	$[-.093; .139]$	Non-significant
	AI	-.015	$[-.113; .074]$	Non-significant
Ack → WARM → PSI → PI	Human	.264	[.165; .383]	Significant
	AI	.098	[.052; .158]	Significant
Ack → WARM → PSI → TRST	Human	.124	[.026; .242]	Significant
	AI	.046	[.009; .096]	Significant
Ack → COMP → PSI → PI	Human	.046	[.018; .083]	Significant
	AI	-.037	$[-.069; -.014]$	Significant
IMM: Ack → WARM → TRST	AI – Human	-.177	$[-.361; -.019]$	Significant
IMM: Ack → WARM → PSI → PI	AI – Human	-.166	$[-.258; -.095]$	Significant
IMM: Ack → COMP → PSI → PI	AI – Human	-.084	$[-.150; -.033]$	Significant
IMM: Ack → COMP → TRST	AI – Human	-.141	$[-.377; .079]$	Non-significant

Note: b = unstandardized indirect-effect estimate; CI = percentile bootstrap confidence interval; IMM = index of moderated mediation (AI – Human); Ack = acknowledgment. The H5b IMM estimate is the mean of the bootstrap contrasts.

The moderating role of need for cognition

Under H6a and H6b, whether the effect of error acknowledgment on the outcome variables is moderated by need for cognition (NFC; mean-centered) was tested with regression models. The Acknowledgment $\times$ NFC interaction was not significant for trust ($b = .106$; $SE = .085$; $p = .216$); therefore, H6a is not supported. For purchase intention, however, the interaction was significant and negative ($b = -.187$; $SE = .083$; $p = .025$): the positive effect of error acknowledgment on purchase intention was strongest at low NFC levels (-1 SD) ($b = 2.203$; $p < .001$) and comparatively weaker at high NFC levels ($+1$ SD) ($b = 1.754$; $p < .001$). This finding supports H6b and is consistent with the view that the acknowledgment cue carries more weight among consumers with a lower propensity for elaboration, who are more likely to rely on peripheral cues.

A summary of the hypothesis tests is presented in Table 8.

Table 8: Summary of Hypothesis Testing Results

Hypothesis	Content	Result
H1	The positive effect of error acknowledgment on warmth is stronger for the human than for the AI influencer	Supported
H2	Error acknowledgment decreases the AI influencer's perceived competence	Supported
H3	The positive effect of error acknowledgment on PSI is stronger for the human than for the AI influencer	Supported
H4a	Warmth mediates the effect of error acknowledgment on trust	Supported (weaker for AI)
H4b	Competence mediates the effect of error acknowledgment on trust	Not supported
H4c	Warmth mediates the effect of error acknowledgment on purchase intention	Not supported
H4d	Competence mediates the effect of error acknowledgment on purchase intention	Not supported
H5a	Error acknowledgment $\rightarrow$ warmth $\rightarrow$ PSI $\rightarrow$ purchase intention (serial mediation)	Supported (weaker for AI)
H5b	Error acknowledgment $\rightarrow$ competence $\rightarrow$ PSI $\rightarrow$ purchase intention (serial mediation)	Supported (opposite signs; IMM significant)
H6a	NFC moderates the effect of error acknowledgment on trust	Not supported
H6b	NFC moderates the effect of error acknowledgment on purchase intention	Supported

Results and discussion

This study examined the effect of human and artificial intelligence influencers' responses to informational errors (acknowledgment vs. denial) on consumer perceptions and intentions through a 2×2 experimental designs. The results point to three main conclusions, each of which is discussed below in relation to the literature.

First, acknowledging the error increased warmth, parasocial interaction, trust, and purchase intention for both influencer types; however, these positive effects emerged markedly more strongly for the human influencer than for the AI influencer. As the simple effects analyses showed, the gains produced by acknowledgment in the human condition (warmth $d = 2.19$; parasocial interaction $d = 1.27$; trust $d = 1.84$; purchase intention $d = 1.93$) were roughly two to four times the size of the corresponding gains in the AI condition ($d = 0.72, 0.63, 0.43$, and 0.52 , respectively). This pattern is consistent with crisis communication findings showing that responsibility-taking apologies generate warmth and trust repair (Chung & Lee, 2021; Yang et al., 2025); yet the asymmetry observed in favor of the human actor extends service recovery findings that human apologies are perceived as more sincere than machine apologies (Choi et al., 2021; Hu et al., 2021) to the influencer marketing context. In the present setting, the returns of error acknowledgment were considerably larger when the actor was human.

Second, and most strikingly, the results revealed an asymmetric effect of error acknowledgment on competence perceptions. The interaction on competence was the largest in the study ($\eta_p^2 = .174$) and took the form of a full crossover: while admitting the error was associated with higher perceived competence for the human influencer ($d = 0.99$), it was associated with lower perceived competence for

the AI influencer ($d = -0.84$). Notably, this opposition of effects cancelled out at the level of the main effect of error response on competence ($p = .258$), underscoring that the phenomenon is visible only through the interaction. The crossover pattern is consistent with the algorithm aversion literature showing that people lose trust in erring algorithms faster than in humans (Berger et al., 2021; Dietvorst et al., 2015) and may be explained by two mechanisms proposed in that literature, although neither these mechanisms nor algorithm aversion itself was measured directly in the present study: the perfection standard expected of machines turns an openly owned error into a violation of high accuracy expectations (Han & Ko, 2025; Molina & Sundar, 2024), while attributing the error to the algorithm's own inadequacy renders it a diagnostic signal for future performance (Berger et al., 2021; Toader et al., 2020). The results also offer experimental evidence relevant to the AI transparency debate: the AI influencer's admission of its error increased perceived warmth but decreased perceived competence. Considering that the empirical accumulation of the pratfall effect in marketing contexts is still limited, this study provides original evidence on the boundary conditions of the effect (its emergence for human actors and its reversal for artificial ones).

Third, the mediation results distinguish the direct and serial indirect routes to purchase intention. The indirect effect of error acknowledgment on trust through warmth and the serial effect on purchase intention through warmth and parasocial interaction were significant in both conditions, with smaller effects in the AI condition ($IMM = -.177$ and $-.166$, respectively). By contrast, neither simple indirect effect on purchase intention through warmth or competence was significant. The competence $\rightarrow$ parasocial interaction $\rightarrow$ purchase intention serial effect was positive for the human influencer ($b = .046$) and negative for the AI influencer ($b = -.037$), and their bootstrap contrast excluded zero ($IMM = -.084$; 95% CI $[-.150; -.033]$). This pattern is consistent with the primacy-of-warmth argument in impression formation (Abele et al., 2021; Brambilla et al., 2012; Kolbl et al., 2019), although the measured paths do not establish the proposed psychological mechanisms on their own. The Acknowledgment $\times$ AI interaction for parasocial interaction was not significant after warmth and competence were included in the model ($\beta = .042$; $p = .475$).

Finally, evaluated from the perspective of the Elaboration Likelihood Model, the need for cognition results suggest that error acknowledgment operates as a heuristic cue mainly among consumers with a lower propensity for elaboration: the positive effect of acknowledgment on purchase intention weakened as need for cognition increased (interaction $b = -.187$; $p = .025$), whereas no such moderation was observed for trust. This pattern is consistent with classic findings that individuals low in need for cognition respond more strongly to signal-type cues (Cacioppo et al., 1983; Chaiken, 1980; Inman et al., 1990); consumers high in need for cognition may not regard an error admission alone as a sufficient indicator of quality and may evaluate the recommendation content in greater detail (Petty & Cacioppo, 1986). Because the processing route was inferred from a dispositional measure rather than observed directly, however, these interpretations should be treated with caution. That the moderation emerged only for purchase intention and not for trust may indicate that trust judgments rest on more deliberate evaluations, whereas purchase intention is more sensitive to heuristic cues; this possibility warrants direct testing in future research.

Theoretical contributions

The study makes four contributions to theory. First, it extends the pratfall effect literature from physical or endearing blunders to the context of informational error and responsibility-taking, and it establishes the actor-type-dependent boundary conditions of the effect: the same admission that was associated with higher perceived competence for the human actor was associated with lower perceived competence for the artificial one. Because the error itself was held constant across all four conditions, this reversal can be attributed to the act of acknowledgment rather than to the error, which extends the algorithm aversion account (Berger et al., 2021; Dietvorst et al., 2015) by suggesting that competence penalties toward algorithmic agents may be triggered not only by observing an algorithm err but also by the algorithm's open admission of the error; because algorithm aversion was not measured directly, this interpretation should be verified in future research using direct measures.

Second, the study integrates the Stereotype Content Model, parasocial interaction, and algorithm aversion approaches within a single model and shows how they connect: the effects of the manipulations travelled primarily through warmth and the warmth $\rightarrow$ parasocial interaction serial path, consistent with the primacy-of-warmth argument (Abele et al., 2021; Brambilla et al., 2012; Kolbl et al., 2019) in the influencer marketing context, while the disappearance of the interactive effect of error acknowledgment and influencer type on parasocial interaction once warmth and competence were controlled suggests that parasocial differences between AI and human influencers are largely mechanism-based (Stein et al., 2024). Third, the findings point to the dual nature of AI transparency in

the eyes of consumers: an error admission by an AI influencer was accompanied by both a warmth gain and a competence loss, providing experimental evidence for a debate that has largely been conducted at the conceptual level. Finally, by modeling elaboration propensity through need for cognition, the study identifies a dispositional boundary condition for the effect of error acknowledgment on purchase intention.

Managerial implications

For brands working with human influencers, the findings are unambiguous: promptly and openly taking responsibility in the event of an error yields a strong increase in trust and purchase intention, whereas denial and attributing the error to external factors carry a heavy cost across all consumer responses. For brands using AI influencers, error communication needs to be designed more delicately. Although an AI influencer's explicit acknowledgment of the error still produces better outcomes than denial in terms of warmth, trust, and purchase intention, the loss in competence perceptions must be taken into account. The findings of chatbot recovery research—that a plain apology can make competence limitations salient and that relationship-oriented or solution-oriented messages can outperform a plain apology (Song et al., 2023)—are instructive in this context: for AI influencers, error acknowledgment should be presented together with corrective action and elements that provide competence assurance (e.g., concrete information that the source of the error has been resolved).

The moderation findings add a targeting nuance. Because the effect of error acknowledgment on purchase intention was strongest among consumers low in need for cognition, acknowledgment messages may be particularly relevant in low-involvement, fast-scrolling social media contexts that may encourage greater reliance on heuristic cues; this possibility should be tested directly in future research. When addressing audiences with a higher propensity for elaboration, an admission alone is unlikely to suffice; brands should accompany it with substantive information that can withstand detailed scrutiny of the message content.

Limitations and future research

The limitations of the study should be noted. First, the data were obtained from a single-country sample and a scenario-based experiment; since attitudes toward virtual influencers and disclosure effects are known to vary across cultural contexts (Khalfallah & Keller, 2025; Muniz et al., 2024; Rizzo et al., 2025), testing the effects in real social media interactions and in different cultural contexts will strengthen external validity. Second, a single error type (a minor informational error) and a single sector context were used; considering findings that apology effectiveness varies with error severity (Grover et al., 2019; Xu & Ling, 2026), the severity, type, and product category of the error should be systematically varied in future studies.

Third, several constructs used to interpret the findings were not measured directly. Algorithm aversion served as the theoretical rationale for the competence penalty observed in the AI condition, the perceived honesty or sincerity of the acknowledgment was not assessed, and elaboration was captured through dispositional need for cognition rather than through measures of the processing route actually used. Future studies should measure these constructs directly (e.g., algorithm aversion, perceived sincerity of the response, and message elaboration through thought-listing procedures) to test the proposed explanations more rigorously.

Fourth, because the mediators and dependent variables were collected at the same measurement occasion, the causal ordering in the mediation findings rests on theoretical justification; although the common method bias test results remained within acceptable limits, longitudinal or multi-wave designs could test the mechanisms more strongly. Finally, the establishment of only partial rather than full scalar invariance means that observed composite score comparisons were made. Future research may focus on post-error recovery strategies, the cumulative effect of repeated errors, the moderating role of error severity, and the timing of AI disclosure.

Conclusion

This study asked how consumers respond when human and AI influencers face an informational error, and whether openly acknowledging responsibility or denying the error is the more advantageous communication strategy for each actor. Based on a 2×2 between-subjects scenario experiment with 496 social media users, the answer is twofold. On the one hand, error acknowledgment proved to be the more favorable strategy for both influencer types across warmth, parasocial interaction, trust, and purchase intention. On the other hand, its returns were far from symmetric: the benefits of acknowledgment were considerably larger for the human influencer, whereas for the AI influencer

acknowledgment carried a measurable competence cost, with the same admission that elevated the human actor's perceived competence lowering that of the artificial one.

Taken together, the findings position error acknowledgment as the sound default in influencer error communication while showing that its consequences are filtered through who—or what—the communicator is. For human influencers, admitting a mistake is an investment that pays off across every consumer response examined; for artificial ones, it is a trade-off in which human-like warmth is purchased at the price of machine-like competence. Understanding and managing this trade-off will become increasingly consequential as AI influencers claim a growing share of brand communication.

Peer-review:

Externally peer-reviewed

Conflict of interests:

The author has no conflict of interest to declare.

Grant Support:

The author declared that this study has received no financial support

Ethical Statement:

Ethics committee approval was received for this study from Sivas Cumhuriyet University, Social and Human Sciences Ethics Committee on 25 November 2024 (document no. 494383).

Data Availability Statement:

The datasets generated and/or analysed during the current study are available from the corresponding author on reasonable request.

Use of Artificial Intelligence:

During the preparation of this work, the author used Claude strictly for spell-checking and grammar-editing purposes. After using this tool, the author carefully reviewed and edited the text, and takes full responsibility for the final content of the publication.

References

- Aaker, J., Vohs, K. D., & Mogilner, C. (2010). Nonprofits are seen as warm and for-profits as competent: Firm stereotypes matter. *Journal of Consumer Research*, 37(2), 224–237. <https://doi.org/10.1086/651566>
- Abele, A. E., Ellemers, N., Fiske, S. T., Koch, A., & Yzerbyt, V. (2021). Navigating the social world: Toward an integrated framework for evaluating self, individuals, and groups. *Psychological Review*, 128(2), 290–314. <https://doi.org/10.1037/rev0000262>
- Allal-Chérif, O., Puertas, R., & Carracedo, P. (2024). Intelligent influencer marketing: How AI-powered virtual influencers outperform human influencers. *Technological Forecasting and Social Change*, 200, 123113. <https://doi.org/10.1016/j.techfore.2023.123113>
- Ao, L., Bansal, R., Pruthi, N., & Khaskheli, M. B. (2023). Impact of social media influencers on customer engagement and purchase intention: A meta-analysis. *Sustainability*, 15(3), 2744. <https://doi.org/10.3390/su15032744>
- Ashraf, A., Hameed, I., & Saeed, S. A. (2023). How do social media influencers inspire consumers' purchase decisions? The mediating role of parasocial relationships. *International Journal of Consumer Studies*, 47(4), 1416–1433. <https://doi.org/10.1111/ijcs.12917>

- Auter, P. J. (1992). Psychometric: TV that talks back: An experimental validation of a parasocial interaction scale. *Journal of Broadcasting & Electronic Media*, 36(2), 173–181. <https://doi.org/10.1080/08838159209364165>
- Baek, T. H., Kim, J., & Kim, J. H. (2026). Effect of disclosing AI-generated content on prosocial advertising evaluation. *International Journal of Advertising*, 45(1), 171–192. <https://doi.org/10.1080/02650487.2024.2401319>
- Belanche, D., Casaló, L. V., & Flavián, M. (2024). Human versus virtual influences, a comparative study. *Journal of Business Research*, 173, 114493. <https://doi.org/10.1016/j.jbusres.2023.114493>
- Berger, B., Adam, M., Rühr, A., & Benlian, A. (2021). Watch me improve-algorithm aversion and demonstrating the ability to learn. *Business & Information Systems Engineering*, 63(1), 55–68. <https://doi.org/10.1007/s12599-020-00678-5>
- Bernritter, S. F., Verlegh, P. W., & Smit, E. G. (2016). Why nonprofits are easier to endorse on social media: The roles of warmth and brand symbolism. *Journal of Interactive Marketing*, 33(1), 27–42. <https://doi.org/10.1016/j.intmar.2015.10.002>
- Brambilla, M., Sacchi, S., Rusconi, P., Cherubini, P., & Yzerbyt, V. Y. (2012). You want to give a good impression? Be honest! Moral traits dominate group impression formation. *British Journal of Social Psychology*, 51(1), 149–166. <https://doi.org/10.1111/j.2044-8309.2010.02011.x>
- 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, 103790. <https://doi.org/10.1016/j.jretconser.2024.103790>
- Cacioppo, J. T., & Petty, R. E. (1982). The need for cognition. *Journal of Personality and Social Psychology*, 42(1), 116–131. <https://doi.org/10.1037/0022-3514.42.1.116>
- Cacioppo, J. T., Petty, R. E., & Kao, C. F. (1984). The efficient assessment of need for cognition. *Journal of Personality Assessment*, 48(3), 306–307. https://doi.org/10.1207/s15327752jpa4803_13
- Cacioppo, J. T., Petty, R. E., & Morris, K. J. (1983). Effects of need for cognition on message evaluation, recall, and persuasion. *Journal of Personality and Social Psychology*, 45(4), 805–818. <https://doi.org/10.1037/0022-3514.45.4.805>
- Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-dependent algorithm aversion. *Journal of Marketing Research*, 56(5), 809–825. <https://doi.org/10.1177/0022243719851788>
- Chaiken, S. (1980). Heuristic versus systematic information processing and the use of source versus message cues in persuasion. *Journal of Personality and Social Psychology*, 39(5), 752–766. <https://doi.org/10.1037/0022-3514.39.5.752>
- Cheng, X., Zhang, X., Cohen, J., & Mou, J. (2022). Human vs. AI: Understanding the impact of anthropomorphism on consumer response to chatbots from the perspective of trust and relationship norms. *Information Processing & Management*, 59(3), 102940. <https://doi.org/10.1016/j.ipm.2022.102940>
- Cheung, G. W., & Rensvold, R. B. (2002). Evaluating goodness-of-fit indexes for testing measurement invariance. *Structural Equation Modeling*, 9(2), 233–255. https://doi.org/10.1207/S15328007SEM0902_5
- Choi, S., Mattila, A. S., & Bolton, L. E. (2021). To err is human(-oid): How do consumers react to robot service failure and recovery?. *Journal of Service Research*, 24(3), 354–371. <https://doi.org/10.1177/1094670520978798>
- Chung, S., & Lee, S. (2021). Crisis management and corporate apology: The effects of causal attribution and apology type on publics' cognitive and affective responses. *International Journal of Business Communication*, 58(1), 125–144. <https://doi.org/10.1177/2329488417735646>
- Cocker, H., Mardon, R., & Daunt, K. L. (2021). Social media influencers and transgressive celebrity endorsement in consumption community contexts. *European Journal of Marketing*, 55(7), 1841–1872. <https://doi.org/10.1108/EJM-07-2019-0567>
- Cuddy, A. J., Fiske, S. T., & Glick, P. (2008). Warmth and competence as universal dimensions of social perception: The stereotype content model and the BIAS map. *Advances in Experimental Social Psychology*, 40, 61–149. [https://doi.org/10.1016/S0065-2601\(07\)00002-0](https://doi.org/10.1016/S0065-2601(07)00002-0)

- Dabiran, E., Farivar, S., Wang, F., & Grant, G. (2024). Virtually human: Anthropomorphism in virtual influencer marketing. *Journal of Retailing and Consumer Services*, 79, 103797. <https://doi.org/10.1016/j.jretconser.2024.103797>
- Dedeoğlu, B. B., Bilgihan, A., Ye, B. H., Wang, Y., & Okumus, F. (2021). The role of elaboration likelihood routes in relationships between user-generated content and willingness to pay more. *Tourism Review*, 76(3), 614–638. <https://doi.org/10.1108/TR-01-2019-0013>
- Diamantopoulos, A., Szöcs, I., Florack, A., Kolbl, Ž., & Egger, M. (2021). The bond between country and brand stereotypes: Insights on the role of brand typicality and utilitarian/hedonic nature in enhancing stereotype content transfer. *International Marketing Review*, 38(6), 1143–1165. <https://doi.org/10.1108/IMR-09-2020-0209>
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>
- 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/10.1177/002224379102800305>
- El Hedhli, K., Zourrig, H., Al Khateeb, A., & Alnawas, I. (2023). Stereotyping human-like virtual influencers in retailing: Does warmth prevail over competence?. *Journal of Retailing and Consumer Services*, 75, 103459. <https://doi.org/10.1016/j.jretconser.2023.103459>
- Eriksson, M. (2018). Lessons for crisis communication on social media: A systematic review of what research tells the practice. *International Journal of Strategic Communication*, 12(5), 526–551. <https://doi.org/10.1080/1553118X.2018.1510405>
- Fiske, S. T. (2018). Stereotype content: Warmth and competence endure. *Current Directions in Psychological Science*, 27(2), 67–73. <https://doi.org/10.1177/0963721417738825>
- Fiske, S. T., Cuddy, A. J., & Glick, P. (2007). Universal dimensions of social cognition: Warmth and competence. *Trends in Cognitive Sciences*, 11(2), 77–83. <https://doi.org/10.1016/j.tics.2006.11.005>
- Folkvord, F., Roes, E., & Bevelander, K. (2020). Promoting healthy foods in the new digital era on Instagram: An experimental study on the effect of a popular real versus fictitious fit influencer on brand attitude and purchase intentions. *BMC Public Health*, 20(1), 1677. <https://doi.org/10.1186/s12889-020-09779-y>
- 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>
- Franke, C., Groeppel-Klein, A., & Müller, K. (2023). Consumers' responses to virtual influencers as advertising endorsers: Novel and effective or uncanny and deceiving?. *Journal of Advertising*, 52(4), 523–539. <https://doi.org/10.1080/00913367.2022.2154721>
- Ge, H., Wang, W., Wang, Y., & Tan, R. (2025). How does original equipment manufacturing brand disclosure affect purchase intention? The mediating role of brand competence and brand warmth. *Asia Pacific Journal of Marketing and Logistics*, 37(8), 2205–2227. <https://doi.org/10.1108/APJML-07-2024-0866>
- Gidaković, P., & Žabkar, V. (2022). The formation of consumers' warmth and competence impressions of corporate brands: The role of corporate associations. *European Management Review*, 19(4), 639–653. <https://doi.org/10.1111/emre.12509>
- Grover, S. L., Abid-Dupont, M. A., Manville, C., & Hasel, M. C. (2019). Repairing broken trust between leaders and followers: How violation characteristics temper apologies. *Journal of Business Ethics*, 155(3), 853–870. <https://doi.org/10.1007/s10551-017-3509-3>
- Han, J., & Ko, D. (2025). Trust formation, error impact, and repair in human-AI financial advisory: A dynamic behavioral analysis. *Behavioral Sciences*, 15(10), 1370. <https://doi.org/10.3390/bs15101370>
- 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>

- Hu, L. T., & Bentler, P. M. (1999). Cutoff criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives. *Structural Equation Modeling*, 6(1), 1–55. <https://doi.org/10.1080/10705519909540118>
- Hu, Y., Min, H., & Su, N. (2021). How sincere is an apology? Recovery satisfaction in a robot service failure context. *Journal of Hospitality & Tourism Research*, 45(6), 1022–1043. <https://doi.org/10.1177/10963480211011533>
- Inman, J. J., McAlister, L., & Hoyer, W. D. (1990). Promotion signal: Proxy for a price cut?. *Journal of Consumer Research*, 17(1), 74–81. <https://doi.org/10.1086/208538>
- Joel-Edgar, S., Chowdhury, S., Nagy, P., & Ren, S. (2025). Virtual influencers in social media versus the metaverse: Mind perception, blame judgements and brand trust. *Journal of Business Research*, 189, 115139. <https://doi.org/10.1016/j.jbusres.2024.115139>
- Khalfallah, D., & Keller, V. (2025). Authenticity, ethics, and transparency in virtual influencer marketing: A cross-cultural analysis of consumer trust and engagement: A systematic literature review. *Acta Psychologica*, 260, 105573. <https://doi.org/10.1016/j.actpsy.2025.105573>
- Kolbl, Ž., Arslanagic-Kalajdzic, M., & Diamantopoulos, A. (2019). Stereotyping global brands: Is warmth more important than competence?. *Journal of Business Research*, 104, 614–621. <https://doi.org/10.1016/j.jbusres.2018.12.060>
- Lombard, M., & Xu, K. (2021). Social responses to media technologies in the 21st century: The media are social actors paradigm. *Human-Machine Communication*, 2, 29–55. <https://doi.org/10.30658/hmc.2.2>
- Lou, C., & Yuan, S. (2019). Influencer marketing: How message value and credibility affect consumer trust of branded content on social media. *Journal of Interactive Advertising*, 19(1), 58–73. <https://doi.org/10.1080/15252019.2018.1533501>
- Lou, C., Kiew, S. T. J., Chen, T., Lee, T. Y. M., Ong, J. E. C., & Phua, Z. (2023). Authentically fake? How consumers respond to the influence of virtual influencers. *Journal of Advertising*, 52(4), 540–557. <https://doi.org/10.1080/00913367.2022.2149641>
- Mahmood, A., & Huang, C. M. (2024). Gender biases in error mitigation by voice assistants. *Proceedings of the ACM on Human-Computer Interaction*, 8(CSCW1), 1–27. <https://doi.org/10.1145/3637337>
- Masuda, H., Han, S. H., & Lee, J. (2022). Impacts of influencer attributes on purchase intentions in social media influencer marketing: Mediating roles of characterizations. *Technological Forecasting and Social Change*, 174, 121246. <https://doi.org/10.1016/j.techfore.2021.121246>
- Molina, M. D., & Sundar, S. S. (2024). Does distrust in humans predict greater trust in AI? Role of individual differences in user responses to content moderation. *New Media & Society*, 26(6), 3638–3656. <https://doi.org/10.1177/14614448221103534>
- Mrad, M., Ramadan, Z., Tóth, Z., Nasr, L., & Karimi, S. (2025). Virtual influencers versus real connections: Exploring the phenomenon of virtual influencers. *Journal of Advertising*, 54(1), 1–19. <https://doi.org/10.1080/00913367.2024.2393711>
- Muniz, F., Stewart, K., & Magalhães, L. (2024). Are they humans or are they robots? The effect of virtual influencer disclosure on brand trust. *Journal of Consumer Behaviour*, 23(3), 1234–1250. <https://doi.org/10.1002/cb.2271>
- Nielsen, Y. A., Pfattheicher, S., & Keijsers, M. (2022). Prosocial behavior toward machines. *Current Opinion in Psychology*, 43, 260–265. <https://doi.org/10.1016/j.copsyc.2021.08.004>
- Petty, R. E., & Cacioppo, J. T. (1986). The elaboration likelihood model of persuasion. *Advances in Experimental Social Psychology*, 19, 123–205. [https://doi.org/10.1016/S0065-2601\(08\)60214-2](https://doi.org/10.1016/S0065-2601(08)60214-2)
- Podsakoff, P. M., MacKenzie, S. B., Lee, J. Y., & Podsakoff, N. P. (2003). Common method biases in behavioral research: A critical review of the literature and recommended remedies. *Journal of Applied Psychology*, 88(5), 879–903. <https://doi.org/10.1037/0021-9010.88.5.879>
- Rizzo, C., Baima, G., Janovská, K., & Bresciani, S. (2025). Navigating the uncertainty path of virtual influencers: Empirical evidence through a cultural lens. *Technological Forecasting and Social Change*, 210, 123896. <https://doi.org/10.1016/j.techfore.2024.123896>

- Roy, R., & Naidoo, V. (2021). Enhancing chatbot effectiveness: The role of anthropomorphic conversational styles and time orientation. *Journal of Business Research*, 126, 23–34. <https://doi.org/10.1016/j.jbusres.2020.12.051>
- 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>
- Shi, J., & Jiang, Z. (2023). Competence or warmth: Why do consumers pay for green advertising?. *Asia Pacific Journal of Marketing and Logistics*, 35(11), 2834–2857. <https://doi.org/10.1108/APJML-01-2023-0002>
- Sokolova, K., & Kefi, H. (2020). Instagram and YouTube bloggers promote it, why should I buy? How credibility and parasocial interaction influence purchase intentions. *Journal of Retailing and Consumer Services*, 53, 101742. <https://doi.org/10.1016/j.jretconser.2019.01.011>
- Song, M., Zhang, H., Xing, X., & Duan, Y. (2023). Appreciation vs. apology: Research on the influence mechanism of chatbot service recovery based on politeness theory. *Journal of Retailing and Consumer Services*, 73, 103323. <https://doi.org/10.1016/j.jretconser.2023.103323>
- Srinivasan, R., & Sarial-Abi, G. (2021). When algorithms fail: Consumers' responses to brand harm crises caused by algorithm errors. *Journal of Marketing*, 85(5), 74–91. <https://doi.org/10.1177/0022242921997082>
- Stein, J.-P., Breves, P. L., & Anders, N. (2024). Parasocial interactions with real and virtual influencers: The role of perceived similarity and human-likeness. *New Media & Society*, 26(6), 3433–3453. <https://doi.org/10.1177/14614448221102900>
- Sung, B., Im, H., & Duong, V. C. (2023). Task type's effect on attitudes towards voice assistants. *International Journal of Consumer Studies*, 47(5), 1772–1790. <https://doi.org/10.1111/ijcs.12946>
- Thomas, V. L., & Fowler, K. (2021). Close encounters of the AI kind: Use of AI influencers as brand endorsers. *Journal of Advertising*, 50(1), 11–25. <https://doi.org/10.1080/00913367.2020.1810595>
- Toader, D. C., Boca, G., Toader, R., Măcelaru, M., Toader, C., Ighian, D., & Rădulescu, A. T. (2020). The effect of social presence and chatbot errors on trust. *Sustainability*, 12(1), 256. <https://doi.org/10.3390/su12010256>
- Xu, Y., & Ling, I. L. (2026). Restoring trust: Gratitude vs. apology in healthcare service recovery. *Journal of Service Theory and Practice*, 36(7), 46–75. <https://doi.org/10.1108/JSTP-05-2025-0176>
- Yang, M. R., Pathak, S., & Huang, S. Z. (2025). Warmth and competence: An empirical investigation of the dual impact of corporate apologies on repairing brand trust. *Journal of Human, Earth, and Future*, 6(1), 95–114. <https://doi.org/10.28991/HEF-2025-06-01-07>
- Yılmazdoğan, O. C., Doğan, R. Ş., & Altıntaş, E. (2021). The impact of the source credibility of Instagram influencers on travel intention: The mediating role of parasocial interaction. *Journal of Vacation Marketing*, 27(3), 299–313. <https://doi.org/10.1177/1356766721995973>
- Yuan, S., & Lou, C. (2020). How social media influencers foster relationships with followers: The roles of source credibility and fairness in parasocial relationship and product interest. *Journal of Interactive Advertising*, 20(2), 133–147. <https://doi.org/10.1080/15252019.2020.1769514>
- Zeithaml, V. A. (1988). Consumer perceptions of price, quality, and value: A means-end model and synthesis of evidence. *Journal of Marketing*, 52(3), 2–22. <https://doi.org/10.1177/002224298805200302>
- Zerilli, J., Bhatt, U., & Weller, A. (2022). How transparency modulates trust in artificial intelligence. *Patterns*, 3(4), 100455. <https://doi.org/10.1016/j.patter.2022.100455>
- Zhang, Q., & Abdullah, F. (2026). Virtual vs. human influencers: AI-mediated trust transfer and brand attachment among female consumers. *Journal of Theoretical and Applied Electronic Commerce Research*, 21(7), 209. <https://doi.org/10.3390/jtaer21070209>