



Article

AI Review Bots vs. Humans in Handling Negative Reviews: Who Builds More Trust?

Yizhen Wei ^{1,2} , David (Jingjun) Xu ^{2,*} and Kai Li ¹

¹ Business School, Nankai University, Tianjin 300071, China; yizhenwei3-c@my.cityu.edu.hk (Y.W.); likai@nankai.edu.cn (K.L.)

² Department of Information Systems, City University of Hong Kong, Hong Kong, China

* Correspondence: davidxu@cityu.edu.hk

Abstract

As artificial intelligence (AI) increasingly engages in managing consumer interactions on e-commerce platforms, an important question arises: how do public AI-generated replies to negative reviews compare with human replies in fostering consumer trust? This study investigates consumer trust in reply source (AI review bots vs. human agents) across three reply strategies, namely, default, thinking, and feeling, in a public review context. AI review bots are defined as automated systems that publicly respond to consumer reviews. Using a controlled laboratory experiment, we find that when using the default strategy, human replies elicit greater trust than AI replies, mediated by higher perceived authenticity and persuasiveness. Conversely, when using the thinking strategy, AI replies outperform human replies in building consumer trust, as they are perceived as more authentic and persuasive. In the feeling strategy, there is a convergence between AI and human replies. These findings demonstrate that the effectiveness of AI versus human replies depends on the strategy adopted. Theoretically, this study extends research on human–AI interaction by introducing a three-strategy framework to systematically compare AI and human communicators in public review contexts. Practically, the results guide e-commerce sellers and platforms on when and how to deploy AI review bots to effectively manage consumer trust in response to negative reviews.

Keywords: AI review bots; consumer trust; negative reviews; human–AI interaction

1. Introduction

On e-commerce platforms, consumers post a vast number of reviews about products and services, and negative reviews are especially influential. For example, one widely used Amazon review dataset contains over 233 million customer reviews across more than 15 million products [1], many of which contain complaints about product quality, delivery, or service experiences. Negative reviews can reduce potential consumers' purchase intentions and damage sellers' reputation and sales performance [2,3]. Consequently, effective management of negative reviews has become a critical concern for online sellers [4]. A range of reply strategies, including apologies, monetary compensation, and expression of emotional support, to negative reviews can mitigate their adverse effects [5–8]. However, as the number of consumer reviews increases, sellers face challenges in managing negative feedback effectively because manual replies are constrained by delays, high resource demands, and the practical impossibility of individually addressing large volumes of negative reviews.



Academic Editor: Hua Pang

Received: 16 November 2025

Revised: 8 March 2026

Accepted: 16 March 2026

Published: 22 March 2026

Copyright: © 2026 by the authors.

Licensee MDPI, Basel, Switzerland.

This article is an open access article distributed under the terms and conditions of the [Creative Commons Attribution \(CC BY\)](https://creativecommons.org/licenses/by/4.0/) license.

To address these challenges, recent advances in AI offer a potential solution through AI review bots—automated systems that generate replies to consumer reviews in *public* online contexts [9]. These bots provide immediate, scalable, and context-aware responses. As AI-mediated responses become increasingly prevalent across digital platforms—for instance, “Reviewer Robert” on Weibo automatically replies to user posts (<https://m.weibo.cn/u/5762999670>) (accessed on 8 March 2026), demonstrating AI’s ability to participate in *public* discussions and manage consumer feedback. If such public AI review bots can effectively replace human sellers in handling negative reviews, firms could substantially reduce time and labor costs while maintaining active customer engagement. However, prior research on AI and human replies has predominantly focused on *private*, one-to-one interactions [10–12]. In such contexts, human replies are generally favored because consumers expect individualized and personalized communication [13,14]. In contrast, *public* review contexts differ fundamentally by shifting the communication dynamic from dyadic problem-resolution to one-to-many impression management. Specifically, *public* negative reviews are visible to a broad audience [9], and consumers often post them not merely to seek problem resolution but to express dissatisfaction or frustration [15,16]. When evaluating these public replies, observing consumers attribute the underlying motives differently: they are less focused on personalized empathy and more concerned with the objective competence and informational value of the response. In such contexts, the consistent and analytical nature of AI replies may render them more persuasive than human replies [17], thereby challenging prior findings established in *private* interactions. Yet, systematic empirical evidence comparing AI and human replies in managing *public* negative reviews remains scarce.

To address this question, we build on prior work that classifies reply types for both AI and human agents [18,19]. We focus on three parallel reply strategies commonly observed in e-commerce practice: *default*, which refers to standardized, template-like replies, corresponding to mechanical AI and templated human replies; *thinking*, which refers to analytical and data-driven responses, corresponding to thinking AI and rational human replies; *feeling*, empathetic and emotionally attuned responses, corresponding to feeling AI and emotional human replies. These strategies allow us to systematically compare AI and human replies under equivalent communicative conditions while maintaining a consistent research focus on the AI vs. human comparison. Accordingly, we address the following research questions:

RQ1. *Can AI review bots effectively replace human sellers when replying to public negative reviews?*

RQ2. *If so, under which reply strategies (default, thinking, feeling) do AI review bots may outperform human replies?*

To answer these questions, we conducted a controlled laboratory experiment comparing AI and human replies across the three parallel strategies. Our findings reveal that, in the default strategy, human replies generate higher levels of consumer trust than AI replies, consistent with prior evidence from *private* interactions. Under the thinking strategy, AI replies outperform human replies, as consumers perceive AI as more authentic and persuasive in delivering data-driven responses. Under the feeling strategy, AI and human replies show convergence overall.

This study makes several contributions. Theoretically, we advance the literature on AI-generated content by extending its examination from *private* service settings to *public* e-commerce reviews. By proposing a strategy-contingent effectiveness framework, we move beyond a simple AI versus human comparison to provide the first systematic comparisons between AI and human replies under equivalent strategies. Furthermore, we examine perceived authenticity and persuasiveness as key mechanisms driving consumer trust, and

more importantly, identify that their relative influence differs across reply strategies when comparing AI with human replies. Methodologically, we establish a rigorous comparative design that conceptually maps AI capabilities onto parallel human reply typologies (default, thinking, and feeling) to ensure equivalent communicative conditions. Practically, our findings provide actionable insights for platforms and online sellers, suggesting specific, strategy-aligned guidelines on when and how platforms should leverage their data analytics capabilities to design and offer customized AI review bots for effectively managing public negative reviews.

The remainder of the paper is structured as follows. Section 2 reviews the relevant literature, focusing on AI and human replies. Section 3 develops our theoretical framework and hypotheses. Section 4 presents the experimental methodology and results. Section 5 concludes with theoretical and practical implications and limitations.

2. Literature Review

2.1. AI vs. Human Replies: From Private to Public Contexts

In *private* reply context, AI and human agents have been widely studied for their abilities to manage customer inquiries, complaints, and requests in one-to-one settings [10–12]. For AI agents, such as customer service chatbots [13] and automated financial assistants [20], prior studies highlight a distinct trade-off between functional efficiency and experiential limitations. While AI enhances efficiency and immediacy in service delivery [21,22], its conversational quality often falls short when interactions lack contextual adaptability [23,24]. Moreover, although users value quick responses, the inclusion of affective cues in AI communication triggers expectation–disconfirmation effects, reducing users’ satisfaction [13]. These findings indicate that AI replies in private contexts are efficient yet remain constrained in personalization. Conversely, the literature on human agents traditionally emphasizes their superiority in emotional competence and contextual flexibility [25,26]. Rather than simply providing prompt redress [12,27], human employees can flexibly integrate emotional and cognitive cues in real time to accurately recognize and regulate customer emotions [28–30]. While AI systems prioritize operational speed, human agents are uniquely equipped to deliver the personalized and context-sensitive service recovery that customers actively seek in private exchanges [13,31–33]. However, these established findings regarding human advantage fail to generalize beyond private settings. In contrast to private exchanges, in public contexts, one-to-many interactions are visible to a broad audience, and customers’ goals shift from personalized problem resolution to public self-expression and impression management [9]. Emerging research has begun to examine the impact of public AI-generated content on user behavior [34]. For instance, some study explores how “Reviewer Robert,” an AI bot on Weibo, generates public replies that stimulate user engagement [9]. Similar practices can be observed on commercial platforms (e.g., SMZDM.com (<https://finance.eastmoney.com/a/202403083006492496.html> (accessed on 8 March 2026))), where AI replies to users in the review section. For public human replies, a rich literature has examined, particularly sellers’ responses to negative online reviews [2,35,36]. Specifically, reply strategies including apologies, monetary compensation, and expressions of emotional support can mitigate consumers’ dissatisfaction [1,5–7]. Yet, despite the growing prevalence of AI in public communication, direct comparisons between the analytical objectivity of AI and the emotional superiority of human replies in public review contexts remain scarce.

To examine whether the relative advantage of human replies observed in private contexts persists in public review settings, *and to investigate under which communication strategies AI’s analytical objectivity might outperform humans’ emotional empathy*, we establish a conceptual mapping between AI and human reply typologies to enable systematic cross-

source comparison. Prior research has conceptualized AI capabilities into three categories: mechanical, thinking, and feeling AI [18]. Specifically, mechanical AI systems possess minimal learning ability and perform repetitive, standardized tasks [18]. Thinking AI engages in logical, analytical processing [37], whereas feeling AI recognizes and responds to human emotions to foster emotional resonance with users [13]. While foundational, existing literature predominantly treats these dimensions as a conceptual typology of intelligences isolated to AI systems (e.g., service robots). Our study advances this stream by operationalizing these conceptual AI capabilities into concrete, empirically testable communication strategies. Analogously, communication and service recovery literature has identified three corresponding human reply types: templated, rational, and emotional [19,38]. Templated human replies are standardized, repetitive, and predefined [39], rational replies are analytical and reason-based [40,41], and emotional replies emphasize empathy and affective connection [42]. Building on this theoretical correspondence, we classify reply strategies into three conceptually aligned categories reflecting comparable communicative complexity. Specifically, we adopt the term “default strategy” rather than “mechanical” to ensure conceptual symmetry. While the term “mechanical” carries an inherent machine bias, making it less suitable for human agents, and “templated” implies manual human operation, the “default” strategy represents a neutral, symmetrical label. Defined as standardized and template-like [43], the default concept is commonly used in e-commerce platform practices to describe pre-set or automated responses (<https://help.manychat.com/hc/en-us/articles/14281159586588-Default-Reply-in-Manychat> (accessed on 8 March 2026)) (e.g., default reviews), making it highly applicable to both mechanical AI replies and templated human replies. The *thinking strategy* captures analytical, data-driven, and logical replies, while the *feeling strategy* represents empathetic communication aimed at emotional repair. Following Huang and Rust [18], we use the terms default, thinking, and feeling strategies to maintain symmetry across AI- and human-generated replies.

Crucially, by integrating human agents into a typology previously reserved for AI, we construct a unified comparative lens. This approach transitions the literature from evaluating AI in isolation to investigating the strategy-agent fit under equivalent strategic conditions. Consequently, it allows us to empirically investigate the role specialization of humans and AI in public contexts, moving the debate beyond a binary “who is better” to a structured exploration of how specific strategic contexts influence the effectiveness of each communicator. Figure 1 summarizes our conceptual development.

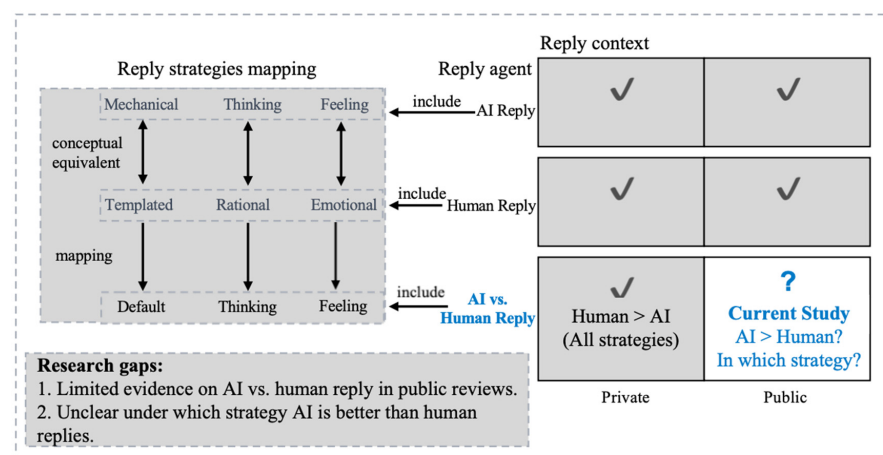


Figure 1. Conceptual development. Note: ✓ indicates relationships or areas that have been examined in prior research.

2.2. Perceived Authenticity and Persuasiveness

Perceived authenticity and persuasiveness are key psychological mechanisms influencing consumer trust and subsequent decision-making behaviors [44,45]. While authenticity is typically conceptualized across historical and categorical dimensions focusing on the origins or criteria of physical entities [46], literature on online reviews primarily emphasizes value authenticity—that is, whether the content reflects the genuine and sincere thoughts of the communicator [47,48]. Rather than simply processing factual information, consumers rely on these expressions of sincerity to evaluate credibility and establish trust [49]. This reliance becomes particularly salient in managing negative reviews, where the perceived authenticity of responses determines their effectiveness in repairing brand reputation [50,51]. Nevertheless, a systematic understanding of how perceived authenticity functions as a mediating mechanism between AI and human replies to build consumer trust under different strategies remains unexplored.

Perceived persuasiveness, defined as the degree to which recipients believe that messages can alter their attitudes or behavioral intentions [52], is conceptually related to perceived authenticity [53]. The elaboration likelihood model (ELM) is typically used in existing studies to examine persuasion in online contexts, positing that persuasion operates via two distinct routes: the central route, based on logical reasoning and evidence; and the peripheral route, based on emotional cues or superficial features [54]. Studies focusing on e-commerce review responses have found that different persuasion strategies, such as attribution of responsibility, explicit promises for rectification, and emotional reassurance [8,55], can significantly shape consumer behavioral intentions. Detailed explanations, combined with forward-looking commitments, are found more persuasive than simple apologies [5]. Moreover, language style, tonal consistency, and logical coherence significantly influence users' assessments of persuasiveness [52]. The interaction between authenticity and persuasiveness has likewise been examined in previous research. When information is perceived as authentic, it is more likely to exert persuasive effects [53,56], thereby enhancing trust [57]. Nevertheless, this mechanism has yet to be systematically tested within the specific context of online review responses, especially when considering the comparison between AI and replies within each of the different strategies, namely, default, thinking, and feeling.

In summary, although the extant literature has underscored the role of perceived authenticity and persuasiveness in shaping consumer trust, how these mechanisms operate between AI and human replies to negative reviews within each of the distinct strategies remains insufficiently understood. In socially interactive environments, such as public response threads on e-commerce platforms, understanding how consumers interpret authenticity, and persuasiveness represents a critical gap.

3. Theoretical Development and Hypotheses

When sellers reply to negative reviews, their primary goal is to mitigate reputational damage and rebuild consumer trust [58–60]. Potential consumers, in turn, infer the underlying motives of such replies, whether they stem from genuine concern or are merely a procedural obligation [4]. According to attribution theory [61], individuals evaluate others' behaviors by attributing them to either intrinsic or extrinsic motivations. These attributions shape consumers' perceptions of the reply's authenticity and persuasiveness, ultimately influencing trust. Building on this framework, we propose that the relative effectiveness of AI vs. human replies depends on the reply strategy employed, namely, default, thinking, and feeling.

The default strategy refers to standardized, template-like replies that lack salient analytical or emotional cues [43]. In such cases, human replies are more likely to be attributed to intrinsic motivations, such as responsibility or concern for customer relation-

ships [19], because even standardized human responses imply a minimal level of effort and engagement [39], indicating intentional concern rather than mere automation. Specifically, potential consumers may infer that human sellers intentionally crafted these replies to protect their reputation or maintain customer satisfaction [62]. In contrast, AI replies following a default strategy are perceived as automated [9,20], stemming from system-level routines rather than deliberate engagement. Consequently, AI replies under this strategy may have lower trust compared to human replies. According to the persuasion knowledge model [63], consumers evaluate the sincerity of persuasive messages by inferring their authenticity. Replies perceived as genuine rather than automated are seen as more persuasive, which in turn enhances trust. Therefore, the effect of reply source (AI vs. human) on trust should be mediated by perceived authenticity and persuasiveness.

Accordingly, we propose the following hypotheses:

H1. *When using the default strategy, human replies generate higher consumer trust than AI replies when replying to negative reviews.*

H2. *When using the default strategy, human replies are perceived as more authentic and persuasive than AI replies, which increases consumer trust.*

The thinking strategy emphasizes analytical, data-driven, and logical reasoning. While human replies using this strategy may still be attributed to intrinsic motivations such as professional diligence [19], AI replies may achieve a stronger fit with consumers' expectations of machine competence. Specifically, the machine heuristic model [64] states that individuals often assume that AI agents are more objective, consistent, and rational than humans. When AI review bots display competence through logic and data analysis, their perceived trustworthiness may exceed that of human agents. Moreover, the thinking strategy aligns with consumers' expectations of AI as rational problem-solvers, thereby reinforcing the perceived trustworthiness of AI-generated replies. Additionally, when using the thinking strategy, human replies adopting highly analytical tones may appear strategic or even insincere, so the logical rigor of AI replies may be perceived as a more reliable indicator of authenticity than human replies.

Therefore, we propose the following hypotheses:

H3. *When using the thinking strategy, AI replies generate higher consumer trust than human replies when replying to negative reviews.*

H4. *When using the thinking strategy, AI replies are perceived as more authentic and persuasive than human replies when replying to negative reviews, which increases consumer trust.*

The feeling strategy emphasizes emotional and empathy expressions. With recent advances in natural language processing and affective computing, AI systems have become increasingly capable of recognizing and reproducing emotional cues [65,66]. When AI review bots express empathy appropriately, consumers may attribute human-like warmth and social presence to them [67]. However, human communicators have been considered superior in emotional expression because they can convey genuine empathy and social understanding [14]. Therefore, because emotional expression remains a core element of human authenticity, we expect that human replies using a feeling strategy should still be perceived as more authentic and persuasive than AI replies. Accordingly, we propose the following hypotheses:

H5. *When using the feeling strategy, human replies generate higher consumer trust than AI replies when replying to negative reviews.*

H6. *When using the feeling strategy, human replies are perceived as more authentic and persuasive than AI replies when replying to negative reviews, which increases consumer trust.*

Figure 2 summarizes our theoretical framework, which focuses on how AI and human replies differentially influence consumer trust in public review contexts. We test these effects across three strategies: default, thinking, and feeling.

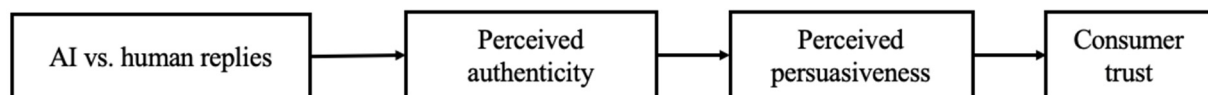


Figure 2. Research framework. Note: The model is tested across three parallel strategies (default, thinking, and feeling).

4. Methods

4.1. Stimulus Materials

Participants were presented with a negative review of a game console, adapted from real consumer reviews on major Chinese e-commerce platforms (e.g., Taobao). A game console was selected as the focal product because it typifies an experience-oriented good whose quality depends on actual usage and is difficult to fully evaluate prior to purchase [68]. Due to this pre-purchase uncertainty, consumers rely heavily on public reviews and subsequent agent replies to assess product reliability and form their trust. Particularly, because negative information is often perceived as more diagnostic than positive reviews in such uncertain contexts [69], prior research highlights that in such experience-based contexts, negative reviews have significant adverse effects on consumer perceptions [70]. This makes the game console contextually appropriate for examining how AI versus human replies can effectively manage public negative reviews and rebuild consumer trust.

Importantly, our study focuses on public review sections to capture the dynamics of one-to-many communication. Unlike private service chats, these public-facing environments are visible to a broad audience of potential consumers, introducing unique effects. Consistent with this setting, the negative review was adapted from actual consumer complaints on Taobao to ensure realism and contextual relevance. Specifically, it highlighted users' dissatisfaction with the product's performance: "A bit disappointed. This console doesn't run smoothly, and the games frequently lag. I regret buying it."

To manipulate the identity of the reply source, we explicitly labeled the agent as either a human seller ("Seller") or an AI review bot ("Reviewer Robert"). The name "Reviewer Robert" was chosen to reflect actual practices on Chinese social media platforms such as Weibo [9]. To maintain internal validity, the content of each reply was identical across the two source conditions within each strategy. Only the source label differed, allowing us to isolate the effect of perceived reply identity (AI vs. human) independent of message content or tone.

To operationalize the three communication strategies, we developed a set of reply scripts based on real review interactions to ensure experimental realism and reliability. Three scripts were created to represent distinct communication strategies: default, thinking, and feeling. The default strategy reflected the standardized, template-like [18] tone commonly used in seller replies on e-commerce platforms: "Thank you for your feedback. A small number of users have reported similar compatibility issues, while the majority found it to be compatible. We will conduct a technical inspection based on your report." The thinking strategy incorporated analytical reasoning and quantitative evidence to signal ana-

lytical competence [37]: “Thank you for your feedback. Upon analysis: (1) Low occurrence rate—only 2.1% of users reported similar compatibility issues; (2) Stable performance—91% of users found it to be compatible; (3) Resolution—we will conduct a technical inspection based on your report.” The feeling strategy emphasized empathy and emotional resonance, adopting affective cues to convey understanding and care [13]: “We totally understand how you feel—if the same happened to me, I’d be disappointed too. Thank you for your feedback. A small number of users have reported similar compatibility issues, while the majority found it to be compatible. We will conduct a technical inspection based on your report.”

The full scripts for all conditions are provided in Appendix A.

4.2. Procedure

We recruited 300 participants from Credamo, a reputable Chinese online experimental platform, and randomly assigned them to one of six experimental conditions, forming three pairwise comparisons: default AI replies vs. default human replies, thinking AI replies vs. thinking human replies, and feeling AI replies vs. feeling human replies (see Table 1). Participants were required to be at least 18 years old and to have prior online purchasing experience. Cases failing to meet these criteria or failing the attention checks were excluded. All participants provided informed consent before participation. The study simulated an online shopping context for a game console. After completing demographic questions, participants viewed product details (including images, price, specifications), followed by a review section containing one negative review and a corresponding reply (see Appendix A). The experimental interface was designed to closely resemble a real e-commerce review page, including product image layout, seller information, and comment formatting, to enhance validity. After reading the product and review content, participants completed several measures. To ensure measurement reliability and validity, all constructs were evaluated using well-established scales adapted from prior literature. Specifically, we assessed perceived authenticity [71] (e.g., “The seller’s reply is truthful”) to capture source sincerity, and perceived persuasiveness [72] (e.g., “I find the reply from the seller persuasive”) to evaluate message persuasiveness. Finally, consumer trust [73] (e.g., “I believe I will be able to use the product as described in the seller’s reply”) was measured as the key outcome. The inclusion and sequencing of these specific constructs reflect the hierarchical psychological process of service recovery, where source-based authenticity serves as a prerequisite for processing message persuasiveness, which ultimately facilitates the formation of consumer trust.

Table 1. Group assignment for the experiment.

Group	Experimental Condition
Group 1	Default AI replies
Group 2	Default human replies
Group 3	Thinking AI replies
Group 4	Thinking human replies
Group 5	Feeling AI replies
Group 6	Feeling human replies

To check the manipulation of the reply source, participants were asked whether they believed the reply was written by a human or by an AI bot. To verify the effectiveness of the strategy manipulation, participants rated the extent to which the reply demonstrated analytical reasoning and emotional understanding, respectively, compared with the default

reply condition. Furthermore, recognizing that participants' preexisting technological experience can influence their reactions to AI agents [74], we measured and controlled for general familiarity with AI as a covariate. Specifically, participants were asked to report their general familiarity with AI (1 = very unfamiliar, 7 = very familiar). All measurement items used are listed in Appendix B.

4.3. Results

Out of the 300 participants recruited, 287 passed the attention check and were included in the analysis. Randomization checks confirmed no significant demographic differences across the six experimental conditions in age, gender, education, or shopping experience (all $ps > 0.10$; see Table 2), confirming the success of random assignment.

Table 2. Randomness Check.

Variable	1: Default AI Replies (N = 49)	2: Default Human Replies (N = 47)	3: Thinking AI Replies (N = 49)	4: Thinking Human Replies (N = 47)	5: Feeling AI Replies (N = 48)	6: Feeling Human Replies (N = 47)	Test Statistic (p-Value)
Age (Mean, SD)	2.53 (0.71)	2.72 (0.90)	2.71 (0.94)	2.38 (0.71)	2.42 (0.74)	2.51 (0.69)	$F = 1.61$ ($p = 0.16$)
Gender (Mean, SD)	0.33 (0.47)	0.47 (0.50)	0.45 (0.50)	0.51 (0.51)	0.44 (0.50)	0.49 (0.51)	$F = 0.81$ ($p = 0.54$)
Education (Mean, SD)	5.29 (0.54)	5.02 (0.49)	5.12 (0.56)	5.15 (0.55)	5.13 (0.57)	5.13 (0.68)	$F = 1.08$ ($p = 0.37$)
Shopping Experience (Mean, SD)	4.14 (0.89)	4.21 (0.75)	4.18 (0.70)	3.96 (0.86)	4.04 (0.71)	4.15 (0.88)	$F = 0.68$ ($p = 0.64$)

Manipulation checks confirmed that participants correctly identified the reply source ($M_{\text{human}} = 0.20$ vs. $M_{\text{AI}} = 0.99$, SDs = 0.40 and 0.08, $p < 0.001$). For the reply strategies manipulation, participants perceived the thinking replies as significantly more analytical (both $ps < 0.001$) and the feeling replies as more emotionally expressive (both $ps < 0.001$), confirming the success of the manipulations. Detailed descriptive statistics are presented in Table 3.

Table 3. Means and Standard Deviations for Different Groups.

Variables	Default AI	Default Human	Thinking AI	Thinking Human	Feeling AI	Feeling Human
Consumer trust	4.16 (1.03)	5.11 (0.81)	5.72 (0.59)	4.21 (1.06)	4.45 (1.23)	4.18 (1.08)
Manipulation check: perceived analytical ability		5.15 (1.00)		5.81 (1.01)		4.86 (1.33)
Manipulation check: perceived emotional expression		3.86 (1.44)		3.48 (1.41)		4.89 (1.32)

Reliability and validity are satisfactory for all constructs (Table 4). Specifically, Cronbach's α values are above 0.70, and AVEs exceed 0.50, indicating acceptable reliability and convergent validity [75,76]. Discriminant validity is also supported because the square root of each construct's AVE is greater than its correlations with other constructs [77].

To test H1, we conducted an ANOVA with the default AI versus human replies as the independent variable and consumer trust as the dependent variable. Results show that default AI replies significantly reduce consumer trust compared with default human replies ($p < 0.001$). This indicates that in default communication contexts, consumers penalize the inherent automation of AI, whereas even templated human replies signal a baseline of intentional concern. This finding aligns with our theoretical expectation that in the absence of analytical or emotional cues, the mere identity of an AI agent inherently limits its ability to foster trust compared to a human baseline. Hence, H1 is supported.

Table 4. Validity and reliability.

Variables	Cronbach's Alpha	Perceived Authenticity	Perceived Persuasiveness	Consumer Trust
Perceived authenticity	0.885	0.832		
Perceived persuasiveness	0.914	0.588	0.895	
Consumer trust	0.846	0.532	0.768	0.809

Note: Diagonal elements are the square root of the average variance extracted (AVE). These values should exceed the inter-construct correlations for adequate discriminant validity. This condition is satisfied for each construct; Convergent validity is adequate when constructs have an AVE of at least 0.5 (the square root of AVE > 0.7); The reliability indicators measured by Cronbach's alpha are all above 0.7.

To further explore the underlying mechanism (H2), we conducted a serial mediation analysis using PROCESS Model 6 with 5000 bootstrap samples [78]. The results indicate that, compared with the default human replies, the default AI replies significantly decrease perceived authenticity ($\beta = -0.72$, $t = -3.84$, $p < 0.001$). Lowered perceived authenticity significantly decreases perceived persuasiveness ($\beta = 0.69$, $t = 5.84$, $p < 0.001$). Thereafter, lowered perceived persuasiveness decreases consumer trust ($\beta = 0.54$, $t = 8.26$, $p < 0.001$). The indirect effects are significant for consumer trust ($\beta = -0.27$, $SE = 0.08$, 95% CI = $[-0.4534, -0.1223]$). This verifies the proposed mechanism, demonstrating that AI identity impairs trust specifically by diminishing the perceived authenticity and subsequent persuasiveness of the reply. Thus, H2 is supported.

To test H3 and H4, we conducted an ANOVA with thinking AI versus human replies as the independent variable. Results show that the thinking AI replies, compared with the thinking human replies, significantly improve consumer trust ($p < 0.001$). This demonstrates that when a recovery strategy relies on analytical information, an AI identity provides a distinct advantage in building consumer trust over a human identity. Hence, H3 is supported. To examine whether or not the effects of thinking AI versus human replies on consumer trust are mediated by perceived authenticity and persuasiveness, we used PROCESS Model 6 with 5000 bootstrap samples [78]. Results show that the thinking AI replies significantly enhance perceived authenticity compared with the thinking human ($\beta = 1.13$, $t = 7.00$, $p < 0.001$). Perceived authenticity significantly increases perceived persuasiveness ($\beta = 0.57$, $t = 4.77$, $p < 0.001$). In turn, perceived persuasiveness significantly increases consumer trust ($\beta = 0.58$, $t = 8.51$, $p < 0.001$). The indirect effect of thinking AI replies (vs. thinking human replies) on consumer trust ($\beta = 0.38$, $SE = 0.12$, 95% CI = $[0.1450, 0.6105]$) through perceived authenticity and persuasiveness is significant. This confirms that the thinking AI not only enhances its authenticity but also serves as a strong signal of persuasiveness, overriding the traditional human advantage. Therefore, the proposed H4 is supported.

To test H5 and H6, we conducted an ANOVA with feeling AI vs. human replies as the independent variable. Results show that the feeling AI replies, compared with the feeling human replies, do not have significant differences in perceived consumer trust ($p = 0.263$). This result demonstrates a meaningful convergence: AI and human agents are equally effective at maintaining consumer trust when employing feeling strategies. We further test the authenticity of feeling AI and human replies. Results show no significant difference ($M_{\text{human}} = 5.05$ vs. $M_{\text{AI}} = 5.20$; $SDs = 1.10$ and 1.18 , $p = 0.524$). Consistent with the findings for trust, this theoretical convergence can be understood through two plausible explanations. First, as AI's capacity for emotional expression improves, consumers may begin to perceive AI and human communicators as equally capable of conveying empathy when handling negative reviews. Second, an alternative explanation lies in the public nature of the review context: consumers might interpret highly emotional human replies as

strategic impression management rather than genuine concern, thereby discounting the traditional human advantage in authenticity.

To confirm the robustness of our mediation structure, we conducted an alternative analysis specifying authenticity and persuasiveness as parallel mediators (Appendix C). Results consistently showed that authenticity and persuasiveness are serial mediators because authenticity alone did not directly predict trust, further validating our serial mediation logic.

Figure 3 summarizes the main model results.

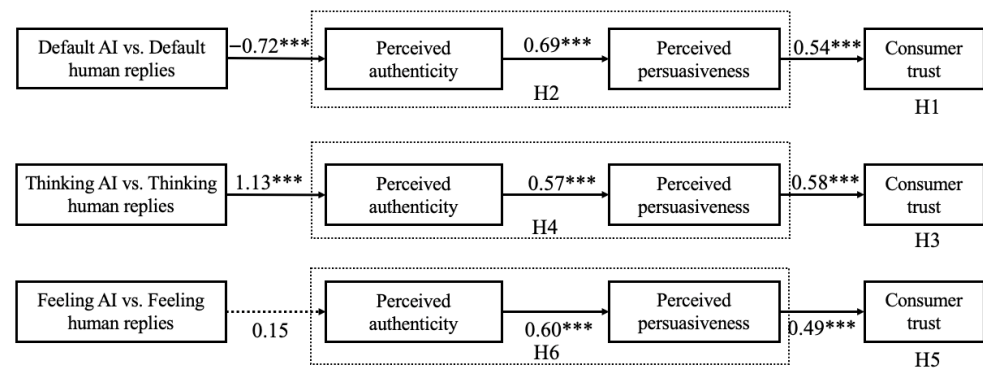


Figure 3. Research results. Note: *** $p < 0.001$.

5. Discussion

5.1. Discussion of Main Findings

This study investigates the comparative effectiveness of AI review bots versus human agents in managing public negative reviews across three communication strategies: default, thinking, and feeling. Moving beyond descriptive comparisons, our findings reveal a profound strategy-agent fit mechanism. In contexts lacking specific strategic cues (default strategy), consumers penalize AI's inherent automation, making human replies more effective for maintaining baseline trust. Conversely, when the context demands analytical rigor and data processing (thinking strategy), AI agents exhibit a distinct advantage, as their perceived objectivity strongly aligns with consumers' expectations of machine competence. Finally, in emotionally driven contexts (feeling strategy), a theoretical convergence occurs, where AI and human agents become functionally equivalent in their capacity to convey empathy. This interplay is consistently mediated by the sequential evaluation of perceived authenticity and persuasiveness.

5.2. Theoretical Implications

This research makes several theoretical contributions. First, this study shifts the theoretical narrative of human–AI interaction from a paradigm of competition to one of role specialization. Prior research, mostly conducted in private one-to-one service contexts [24,67,79], has consistently shown that consumers prefer human responses compared with AI [80,81]. However, we demonstrate that this human superiority does not necessarily hold in public review settings. By systematically comparing AI and human replies, our findings reveal that the effectiveness of an agent is contingent upon the strategy employed, highlighting that the strategy–agent fit matters significantly more than the agent's identity alone. Rather than asking “who is better,” our framework redefines the theoretical boundary by identifying the specific contexts where AI assumes specialized roles (e.g., analytical superiority) or achieves functional equivalence with humans.

Second, we propose and validate a three-strategy comparative framework, namely, default, thinking, and feeling, that serves to unpack the strategy-agent fit. While our typology

draws inspiration from the conceptual classification of AIs (e.g., mechanical, thinking, and feeling AI [18]), our research provides a crucial theoretical advancement. Prior literature predominantly treats these dimensions as conceptual capabilities isolated to AI systems. We advance this stream by operationalizing these conceptual intelligences into empirically testable communication strategies and integrating human agents [19] to construct a unified comparative lens. By doing so, our findings reveal how different strategies activate distinct evaluative mechanisms. Specifically, in the default strategy, the absence of cognitive or affective cues amplifies the perceived automation of AI, whereas even templated human replies signal a baseline of intrinsic concern, reinforcing the traditional human advantage. Conversely, the thinking strategy activates a “machine heuristic”; AI’s data-driven objectivity perfectly aligns with consumer expectations of machine competence, allowing AI to surpass human effectiveness. Crucially, by applying this framework, we identify a theoretically meaningful human-AI convergence in the feeling strategy. On one hand, as AI systems become more sophisticated, their capacity to convey realistic empathy has significantly improved. On the other hand, in public review environments, highly emotional human displays are often discounted by audiences as strategic impression management rather than genuine concern. Furthermore, as consumers experience emotional saturation from encountering repetitive and routinized empathetic responses across e-commerce platforms, the perceived value of human emotional labor diminishes. Consequently, these context-specific mechanisms neutralize the traditional human advantage, rendering AI and human communicators equally viable in emotional recovery.

Third, we advance the understanding of the underlying psychological mechanisms by demonstrating that perceived authenticity and persuasiveness sequentially mediate the effects of AI and human replies, with their functional roles varying across different strategies. While authenticity and persuasiveness have been identified as key constructs in both AI and human communication [38,53,56], prior studies have typically treated them as uniform mediators. Our results reveal a differentiated pattern: in the default strategy, authenticity and persuasiveness amplify human advantage; in the thinking strategy, they enhance AI’s analytical-based ability; and in the feeling strategy, their effects converge, suggesting that advances in affective AI have narrowed the emotional authenticity gap with humans. These mechanisms refine existing theorizing by demonstrating that the same mediators can operate through different psychological routes, depending on the reply type and strategy.

5.3. Practical Implications

Our findings offer actionable guidance for e-commerce sellers and platforms seeking to manage consumer trust when replying to negative reviews. For online sellers, our results show that the effectiveness of replies depends on the strategy employed. When using the default strategy, human replies generate higher levels of consumer trust than AI replies. Therefore, when platforms only provide basic, template-based AI functions, sellers are advised to manually reply to negative reviews rather than relying on default AI replies. Manual replies under this condition are perceived as more authentic and persuasive, making them more effective in mitigating negative consumer perceptions. When using the thinking strategy, AI replies outperform human replies in enhancing consumer trust. Thinking AI replies convey analytical reasoning and data-driven explanations that strengthen perceived authenticity and persuasiveness. Hence, when technically and economically feasible, sellers should prioritize the use of thinking AI review bots for handling negative reviews. When using the feeling strategy, our findings show no significant difference between AI and human replies in perceived authenticity or trust. In this case, sellers can flexibly choose between the two depending on operational constraints. When the number

of negative reviews is large, deploying feeling AI review bots is a cost-effective solution that reduces response time and labor costs. Conversely, when the volume of negative reviews is relatively low, manual replies using the feeling strategy remain a valuable option, as they allow sellers to maintain a more personalized and emotionally attuned connection with consumers.

For platforms, the key implication is to move beyond providing only basic, template-based AI replies (i.e., default AI replies). Although easy to implement, such replies tend to lower perceived authenticity and may undermine consumer trust. Platforms should therefore invest in developing AI review bots with advanced cognitive and emotional capabilities, enabling sellers to utilize both thinking and feeling AI replies as needed. These options not only help sellers align reply strategies with their own operational resources but also allow platforms to differentiate their service offerings. Moreover, because developing thinking and feeling AI functions involves different levels of technical sophistication, platforms may design tiered pricing models that correspond to the complexity and capability of the AI functions provided.

6. Limitations and Future Research

This study has several limitations that open avenues for future research. First, our findings are derived from a controlled laboratory experiment using a single product category and a specific negative review scenario. While this design was essential to maximize internal validity and isolate the causal effects of our communication strategies by strictly controlling for confounding variables (e.g., brand familiarity, review valence), it may limit the external generalizability of the results to real-world e-commerce environments. Future field or platform-based experiments could test whether the observed effects of AI versus human replies on consumer trust remain robust across diverse product categories, varying levels of negative feedback, and more complex, dynamic, and interactive review settings. Particularly, future studies should prioritize theory-driven extensions, such as exploring hybrid AI-human systems where AI handles thinking tasks while human agents intervene for complex feeling strategies, thereby maximizing the synergy between machine competence and human authenticity.

Second, our theoretical framework focuses on authenticity and persuasiveness as key mechanisms explaining how consumers evaluate AI and human replies under different strategies. Although this framework captures how these mechanisms vary across default, thinking, and feeling strategies, the specific alternative drivers of the human-AI convergence in the feeling strategy warrant further empirical investigation. While we theoretically posit that public-context skepticism and emotional saturation neutralize the human advantage, our current study did not directly measure these constructs. Future research could explicitly test these alternative mechanisms by measuring consumers' perceptions of strategic impression management or emotional fatigue in public review sections. This would deepen our understanding of exactly why and when human emotional labor is discounted in digital spaces. Moreover, examining how platforms can implement adaptive strategy switching, dynamically transitioning between AI and human agents based on real-time consumer sentiment analysis, could deepen our understanding of optimal human-AI communication adaptation.

Third, this study focuses on negative consumer reviews, which represent a critical context for reputation recovery and trust management. Future research could extend this framework to positive or neutral review scenarios, where reply strategies may serve different functions, such as reinforcing brand attachment or encouraging advocacy. Ultimately, this research fundamentally shifts the research on human-AI interaction from a simplistic competition of "who is better" to a nuanced, context-dependent understanding of "when

and why” each agent excels. By demonstrating that the strategy–agent fit influences communication effectiveness, our study provides a blueprint for the future of AI-mediated trust in e-commerce. As affective and cognitive AI systems continue to evolve, the future of online service recovery will not be defined by the complete substitution of human labor, but rather by the strategic deployment of AI to foster deeper and more resilient consumer trust.

Author Contributions: Conceptualization, Y.W. and D.X.; Writing—original draft, Y.W.; Writing—review & editing, D.X.; Supervision, D.X. and K.L. All authors have read and agreed to the published version of the manuscript.

Funding: This research was supported by the National Natural Science Foundation of China (Grant Nos. 72572088 and 32400939).

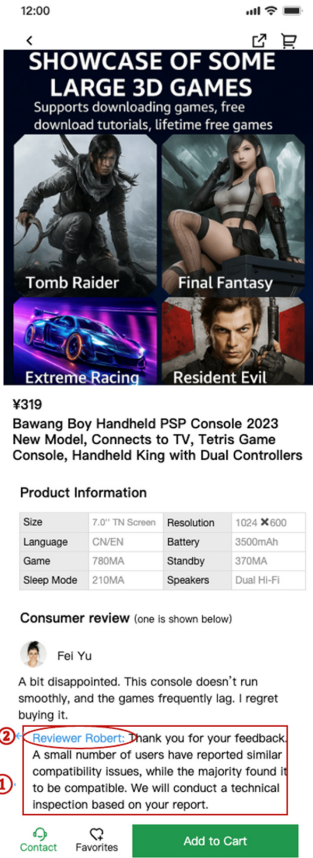
Institutional Review Board Statement: The study was conducted in accordance with the Declaration of Helsinki, and approved by the Institutional Review Board of the Department of Management Science and Engineering, Nankai University (no code attached) on 12 Jan 2025.

Informed Consent Statement: Informed consent was obtained from all subjects involved in the study.

Data Availability Statement: The data presented in this study are available on request from the corresponding author.

Conflicts of Interest: The authors declare no conflict of interest.

Appendix A. Scripts Used in the Experiment

Translated Stimulus Material (The Original Stimulus Materials Were Presented in Chinese)	Description
	① Reply content of default strategy. Thank you for your feedback. <i>A small number of users have reported similar compatibility issues, while the majority found it to be compatible. We will conduct a technical inspection based on your report.</i> ① Reply content of thinking strategy Thank you for your feedback. <i>Upon analysis:</i> (1) <i>Low occurrence rate: only 2.1% of users reported similar compatibility issues.</i> (2) <i>Stable performance: 91% of users found it to be compatible.</i> (3) <i>Resolution: We will conduct a technical inspection based on your report.</i> ① Reply content of feeling strategy <i>We totally understand how you feel—if the same happened to me, I’d be disappointed too.</i> Thank you for your feedback. A small number of users have reported similar compatibility issues, while the majority found it to be compatible. We will conduct a technical inspection based on your report. ② Reply source To differentiate between replies from human sellers and AI review bots, the identity of the reply source is explicitly labeled in the interface as either “Seller” (for human replies) or “Reviewer Robert” (for AI replies). In the AI review bot scenario, participants are told Reviewer Robert is an AI review bot that learns human behavior by mimicking real users on social media and can autonomously interact with others. As an intelligent agent, it can observe and reply in social settings, blending into human networks.

Appendix B. Variables Measured in the Experiments

Constructs	Items
Consumer trust (Wongkitrungrueng and Assarut, 2020) [73]	-I believe the product is as I imagined. -I believe I will be able to use the product as described in the seller/Reviewer Robert's reply. -I believe the product I receive will match the description in the seller/Reviewer Robert's reply.
Perceived authenticity (Yi et al., 2018) [71]	-Reviewer Robert/human seller's reply is truthful. -Reviewer Robert/human seller's reply is false. -Reviewer Robert/human seller's reply is fabricated. -Reviewer Robert/human seller's reply has no reference.
Perceived persuasiveness (Chang et al., 2020) [72]	-I find the reply from the seller/Reviewer Robert unconvincing. -I find the reply from the seller/Reviewer Robert reliable. -I find the reply from the seller/Reviewer Robert persuasive.
Note: Scale ranges from 1 (Very strongly disagree) to 7 (Very Strongly agree).	

Manipulation and Attention Checks	Check Items
Manipulation checks for identification of the reply source	While viewing the reviews, did the page include a reply from an AI bot? -Yes -No
Manipulation checks for thinking strategy (7-point scale)	-In your opinion, to what extent did the reply from the seller/Reviewer Robert express analytical ability?
Manipulation checks for feeling strategy (7-point scale)	-In your opinion, to what extent did the reply from the seller/Reviewer Robert express emotion?

Appendix C. Parallel Mediation

To further validate the robustness of our serial mediation model, we conducted an alternative analysis specifying perceived authenticity and persuasiveness as parallel mediators between reply source and consumer trust within each of the three strategies. Results show that perceived persuasiveness consistently serves as a significant mediator, whereas the direct effect of authenticity on trust is weaker and even non-significant. For example, in the default group, the indirect effect via persuasiveness is significant (95% CI = [−0.9703, −0.3096]) while authenticity was not (95% CI = [−0.1440, 0.1496]). Similar patterns emerge for the thinking group (persuasiveness: 95% CI = [0.4820, 1.3135]; authenticity: 95% CI = [−0.1798, 0.4257]) and feeling group (persuasiveness: 95% CI = [0.0189, 0.6237]; authenticity: 95% CI = [−0.0615, 0.1656]). Figure A1 shows the details. Overall, the results reinforce the serial mediation model, highlighting that authenticity primarily influences persuasiveness, which impacts consumer trust.

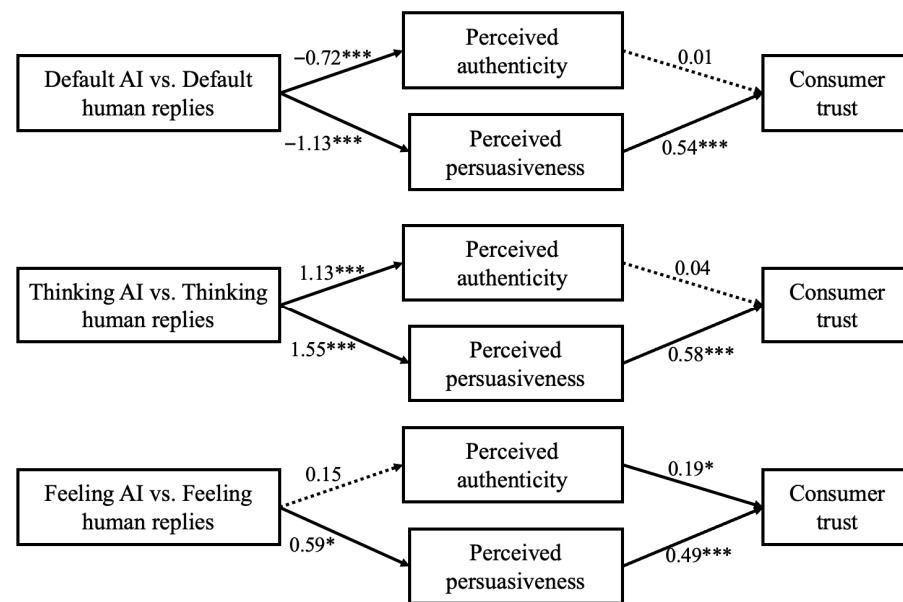


Figure A1. Robustness check. Note: * $p < 0.05$, *** $p < 0.001$.

References

- Wang, F.; Du, Z.; Wang, S. Information multidimensionality in online customer reviews. *J. Bus. Res.* **2023**, *159*, 113727. [\[CrossRef\]](#)
- Esmark Jones, C.L.; Stevens, J.L.; Breazeale, M.; Spaid, B.I. Tell it like it is: The effects of differing responses to negative online reviews. *Psychol. Mark.* **2018**, *35*, 891–901. [\[CrossRef\]](#)
- Utz, S.; Matzat, U.; Snijders, C. On-line reputation systems: The effects of feedback comments and reactions on building and rebuilding trust in on-line auctions. *Int. J. Electron. Commer.* **2009**, *13*, 95–118. [\[CrossRef\]](#)
- Chang, H.H.; Tsai, Y.C.; Wong, K.H.; Wang, J.W.; Cho, F.J. The effects of response strategies and severity of failure on consumer attribution with regard to negative word-of-mouth. *Decis. Support Syst.* **2015**, *71*, 48–61. [\[CrossRef\]](#)
- Wei, C.; Liu, M.W.; Keh, H.T. The road to consumer forgiveness is paved with money or apology? The roles of empathy and power in service recovery. *J. Bus. Res.* **2020**, *118*, 321–334. [\[CrossRef\]](#)
- Hill Cummings, K.; Yule, J.A. Tailoring service recovery messages to consumers' affective states. *Eur. J. Mark.* **2020**, *54*, 1675–1702. [\[CrossRef\]](#)
- Ahmad, F.; Guzmán, F. Negative online reviews, brand equity and emotional contagion. *Eur. J. Mark.* **2021**, *55*, 2825–2870. [\[CrossRef\]](#)
- Wang, K.Y.; Chih, W.H.; Honora, A. How the emoji use in apology messages influences customers' responses in online service recoveries: The moderating role of communication style. *Int. J. Inf. Manag.* **2023**, *69*, 102618. [\[CrossRef\]](#)
- Gao, Y.; Zhang, M.M.; Lysyakov, M. Does social bot help socialize? Evidence from a microblogging platform. *Inf. Syst. Res.* **2025**. [\[CrossRef\]](#)
- Cheng, X.; Zhang, X.; Cohen, J.; Mou, J. Human vs. AI: Understanding the impact of anthropomorphism on consumer response to chatbots from the perspective of trust and relationship norms. *Inf. Process. Manag.* **2022**, *59*, 102940. [\[CrossRef\]](#)
- Pantano, E.; Scarpi, D. I, robot, you, consumer: Measuring artificial intelligence types and their effect on consumers emotions in service. *J. Serv. Res.* **2022**, *25*, 583–600. [\[CrossRef\]](#)
- Davidow, M. Organizational responses to customer complaints: What works and what doesn't. *J. Serv. Res.* **2003**, *5*, 225–250. [\[CrossRef\]](#)
- Han, E.; Yin, D.; Zhang, H. Bots with feelings: Should AI agents express positive emotion in customer service? *Inf. Syst. Res.* **2023**, *34*, 1296–1311. [\[CrossRef\]](#)
- Yang, H.; Xu, H.; Zhang, Y.; Liang, Y.; Lyu, T. Exploring the effect of humor in robot failure. *Ann. Tour. Res.* **2022**, *95*, 103425. [\[CrossRef\]](#)
- Yin, Q.; Sadowski, S. A tale of two platforms: A comparative analysis of language use in consumer complaints on Reddit and Spotify Community. *J. Consum. Behav.* **2024**, *23*, 2071–2086. [\[CrossRef\]](#)
- Grewal, L.; Stephen, A.T.; Bart, Y. Online Venting: The Impact of Temporal Proximity Cues and Emotionality on Perceptions of Negative Online Review Value. *J. Mark. Res.* **2025**, *63*, 27–46. [\[CrossRef\]](#)
- Garvey, A.M.; Kim, T.; Duhachek, A. Bad news? Send an AI. Good news? Send a human. *J. Mark.* **2023**, *87*, 10–25. [\[CrossRef\]](#)
- Huang, M.H.; Rust, R.T. Engaged to a robot? The role of AI in service. *J. Serv. Res.* **2021**, *24*, 30–41. [\[CrossRef\]](#)

19. Ravichandran, T.; Deng, C. Effects of managerial response to negative reviews on future review valence and complaints. *Inf. Syst. Res.* **2023**, *34*, 319–341. [\[CrossRef\]](#)
20. Cheng, X.; Qiao, L.; Yang, B.; Li, Z. An investigation on the influencing factors of elderly people's intention to use financial AI customer service. *Internet Res.* **2024**, *34*, 690–717. [\[CrossRef\]](#)
21. Perez-Vega, R.; Kaartemo, V.; Lages, C.R.; Razavi, N.B.; Männistö, J. Reshaping the contexts of online customer engagement behavior via artificial intelligence: A conceptual framework. *J. Bus. Res.* **2021**, *129*, 902–910. [\[CrossRef\]](#)
22. Jiang, H.; Cheng, Y.; Yang, J.; Gao, S. AI-powered chatbot communication with customers: Dialogic interactions, satisfaction, engagement, and customer behavior. *Comput. Hum. Behav.* **2022**, *134*, 107329. [\[CrossRef\]](#)
23. Janssen, A.; Rodríguez Cardona, D.; Breitner, M.H. More than FAQ! Chatbot taxonomy for business-to-business customer services. In *International Workshop on Chatbot Research and Design*; Springer International Publishing: Cham, Switzerland, 2020; pp. 175–189.
24. Hsu, C.L.; Lin, J.C.C. Understanding the user satisfaction and loyalty of customer service chatbots. *J. Retail. Consum. Serv.* **2023**, *71*, 103211. [\[CrossRef\]](#)
25. Kelley, S.W.; Davis, M.A. Antecedents to customer expectations for service recovery. *J. Acad. Mark. Sci.* **1994**, *22*, 52–61. [\[CrossRef\]](#)
26. Tax, S.S.; Brown, S.W.; Chandrashekar, M. Customer evaluations of service complaint experiences: Implications for relationship marketing. *J. Mark.* **1998**, *62*, 60–76. [\[CrossRef\]](#)
27. Wirtz, J.; Mattila, A.S. Consumer responses to compensation, speed of recovery and apology after a service failure. *Int. J. Serv. Ind. Manag.* **2004**, *15*, 150–166. [\[CrossRef\]](#)
28. Delcourt, C.; Gremler, D.D.; van Riel, A.C.; Van Birgelen, M.J. Employee emotional competence: Construct conceptualization and validation of a customer-based measure. *J. Serv. Res.* **2016**, *19*, 72–87. [\[CrossRef\]](#)
29. Fernandes, T.; Morgado, M.; Rodrigues, M.A. The role of employee emotional competence in service recovery encounters. *J. Serv. Mark.* **2018**, *32*, 835–849. [\[CrossRef\]](#)
30. Maxham, J.G., III; Netemeyer, R.G. Modeling customer perceptions of complaint handling over time: The effects of perceived justice on satisfaction and intent. *J. Retail.* **2002**, *78*, 239–252. [\[CrossRef\]](#)
31. Bock, D.E.; Mangus, S.M.; Folse, J.A.G. The road to customer loyalty paved with service customization. *J. Bus. Res.* **2016**, *69*, 3923–3932. [\[CrossRef\]](#)
32. Lee, E.J.; Park, J.K. Online service personalization for apparel shopping. *J. Retail. Consum. Serv.* **2009**, *16*, 83–91. [\[CrossRef\]](#)
33. Gwinner, K.P.; Bitner, M.J.; Brown, S.W.; Kumar, A. Service customization through employee adaptiveness. *J. Serv. Res.* **2005**, *8*, 131–148. [\[CrossRef\]](#)
34. Bian, Z.; Che, C. How AI Overview of Customer Reviews Influences Consumer Perceptions in E-Commerce? *J. Theor. Appl. Electron. Commer. Res.* **2025**, *20*, 315. [\[CrossRef\]](#)
35. Alba, G.; Slongo, L.A. Getting a no-reply is also a reply: An investigation of unreplied consumer attributions. *J. Retail. Consum. Serv.* **2020**, *54*, 102054. [\[CrossRef\]](#)
36. Crijns, H.; Cauberghe, V.; Hudders, L.; Claeys, A.S. How to deal with online consumer comments during a crisis? The impact of personalized organizational responses on organizational reputation. *Comput. Hum. Behav.* **2017**, *75*, 619–631. [\[CrossRef\]](#)
37. Kyung, N.; Kwon, H.E. Rationally trust, but emotionally? The roles of cognitive and affective trust in laypeople's acceptance of AI for preventive care operations. *Prod. Oper. Manag.* **2025**. *OnlineFirst*. [\[CrossRef\]](#)
38. Zhang, J.; Lu, J.; Wang, X.; Liu, L.; Feng, Y. Emotional expressions of care and concern by customer service chatbots: Improved customer attitudes despite perceived inauthenticity. *Decis. Support Syst.* **2024**, *186*, 114314. [\[CrossRef\]](#)
39. Ding, Z.; Zhang, Y.; Sun, J.; Goh, M.; Yang, Z. To reveal or conceal: AI identity disclosure strategies for merchants. *Int. J. Prod. Econ.* **2025**, *283*, 109564. [\[CrossRef\]](#)
40. Ku, C.H.; Chang, Y.C.; Wang, Y. How to strategically respond to online hotel reviews: A strategy-aware deep learning approach. *Inf. Manag.* **2024**, *61*, 103970. [\[CrossRef\]](#)
41. Wang, W.; Qiu, L.; Kim, D.; Benbasat, I. Effects of rational and social appeals of online recommendation agents on cognition-and affect-based trust. *Decis. Support Syst.* **2016**, *86*, 48–60. [\[CrossRef\]](#)
42. Wan, H.; Mei, M.Q.; Yan, J.; Xiong, J.; Wang, L. How does apology matter? Responding to negative customer reviews on online-to-offline platforms. *Electron. Commer. Res. Appl.* **2023**, *61*, 101291. [\[CrossRef\]](#)
43. An, H.; Li, W.; Yu, Y.; Wang, Z. Biases in online reviews: The default positive review rule and the conditional rebate strategy. *Internet Res.* **2025**, *ahead-of-print*. [\[CrossRef\]](#)
44. Chen, X.; Hyun, S.S.; Lee, T.J. The effects of parasocial interaction, authenticity, and self-congruity on the formation of consumer trust in online travel agencies. *Int. J. Tour. Res.* **2022**, *24*, 563–576. [\[CrossRef\]](#)
45. Touré-Tillery, M.; McGill, A.L. Who or what to believe: Trust and the differential persuasiveness of human and anthropomorphized messengers. *J. Mark.* **2015**, *79*, 94–110. [\[CrossRef\]](#)
46. Newman, G.E. The psychology of authenticity. *Rev. Gen. Psychol.* **2019**, *23*, 8–18. [\[CrossRef\]](#)
47. Zhang, Z.; Patrick, V.M. Mickey D's has more street cred than McDonald's: Consumer brand nickname use signals information authenticity. *J. Mark.* **2021**, *85*, 58–73. [\[CrossRef\]](#)

48. Valsesia, F.; Diehl, K. Let me show you what I did versus what I have: Sharing experiential versus material purchases alters authenticity and liking of social media users. *J. Consum. Res.* **2022**, *49*, 430–449. [[CrossRef](#)]
49. Kim, M.; Kim, J. The influence of authenticity of online reviews on trust formation among travelers. *J. Travel Res.* **2020**, *59*, 763–776. [[CrossRef](#)]
50. Cinelli, M.D.; LeBoeuf, R.A. Keeping it real: How perceived brand authenticity affects product perceptions. *J. Consum. Psychol.* **2020**, *30*, 40–59. [[CrossRef](#)]
51. Lin, W.C.; Lu, T.E.; Peng, M.Y. Service failure recovery on customer recovery satisfaction for airline industry: The moderator of brand authenticity and perceived authenticity. *Manag. Decis. Econ.* **2021**, *42*, 1079–1088. [[CrossRef](#)]
52. Lei, Z.; Yin, D.; Zhang, H. Focus within or on others: The impact of reviewers' attentional focus on review helpfulness. *Inf. Syst. Res.* **2021**, *32*, 801–819. [[CrossRef](#)]
53. Li, Y.; Hou, R.; Tan, R. How customers respond to chatbot anthropomorphism: The mediating roles of perceived humanness and perceived persuasiveness. *Eur. J. Mark.* **2024**, *58*, 2757–2790. [[CrossRef](#)]
54. Cyr, D.; Head, M.; Lim, E.; Stibe, A. Using the elaboration likelihood model to examine online persuasion through website design. *Inf. Manag.* **2018**, *55*, 807–821. [[CrossRef](#)]
55. Treviño, T.; Castaño, R. How should managers respond? Exploring the effects of different responses to negative online reviews. *Int. J. Leis. Tour. Mark.* **2013**, *3*, 237–251. [[CrossRef](#)]
56. Li, X.; Zhu, D.H.; Chang, Y. The negative effect of name: Mentions of frontline service employee name reduce online review persuasiveness. *J. Serv. Res.* **2024**, *27*, 636–652. [[CrossRef](#)]
57. Xu, Y.; Jeong, E.L. Persuasive strategies for encouraging food waste reduction: A study of temporal focus, message framing, and the role of perceived trustworthiness in restaurant. *J. Hosp. Tour. Manag.* **2024**, *60*, 239–251. [[CrossRef](#)]
58. Kumar, N.; Qiu, L.; Kumar, S. Exit, voice, and response on digital platforms: An empirical investigation of online management response strategies. *Inf. Syst. Res.* **2018**, *29*, 849–870. [[CrossRef](#)]
59. Wang, H.; Du, R.; Shen, W.; Qiu, L.; Fan, W. Product reviews: A benefit, a burden, or a trifle? How seller reputation affects the role of product reviews. *MIS Q.* **2022**, *46*, 1243–1272. [[CrossRef](#)]
60. Yang, M.; Zheng, Z.; Mookerjee, V.; Chen, H. Responding to online reviews in competitive markets: A controlled diffusion approach. *MIS Q.* **2023**, *47*, 161–194. [[CrossRef](#)]
61. Kelley, H.H.; Michela, J.L. Attribution theory and research. *Annu. Rev. Psychol.* **1980**, *31*, 457–501. [[CrossRef](#)]
62. Lee, B.K. Hong Kong consumers' evaluation in an airline crash: A path model analysis. *J. Public Relat. Res.* **2005**, *17*, 363–391. [[CrossRef](#)]
63. Ham, C.D.; Nelson, M.R.; Das, S. How to measure persuasion knowledge. *Int. J. Advert.* **2015**, *34*, 17–53. [[CrossRef](#)]
64. Wu, X.; Li, S.; Guo, Y.; Fang, S. Human or AI robot? Who is fairer on the service organizational frontline. *J. Bus. Res.* **2024**, *181*, 114730. [[CrossRef](#)]
65. Lv, X.; Yang, Y.; Qin, D.; Cao, X.; Xu, H. Artificial intelligence service recovery: The role of empathic response in hospitality customers' continuous usage intention. *Comput. Hum. Behav.* **2022**, *126*, 106993. [[CrossRef](#)]
66. Yun, J.; Park, J. The effects of chatbot service recovery with emotion words on customer satisfaction, repurchase intention, and positive word-of-mouth. *Front. Psychol.* **2022**, *13*, 922503. [[CrossRef](#)]
67. Schanke, S.; Burtch, G.; Ray, G. Estimating the impact of "humanizing" customer service chatbots. *Inf. Syst. Res.* **2021**, *32*, 736–751. [[CrossRef](#)]
68. Philp, M.; Nepomuceno, M.V. How reviews influence product usage post-purchase: An examination of video game playtime. *J. Bus. Res.* **2024**, *172*, 114456. [[CrossRef](#)]
69. Lee, J.; Park, D.H.; Han, I. The effect of negative online consumer reviews on product attitude: An information processing view. *Electron. Commer. Res. Appl.* **2008**, *7*, 341–352. [[CrossRef](#)]
70. Colmekcioglu, N.; Marvi, R.; Foroudi, P.; Okumus, F. Generation, susceptibility, and response regarding negativity: An in-depth analysis on negative online reviews. *J. Bus. Res.* **2022**, *153*, 235–250. [[CrossRef](#)]
71. Yi, X.; Fu, X.; Yu, L.; Jiang, L. Authenticity and loyalty at heritage sites: The moderation effect of postmodern authenticity. *Tour. Manag.* **2018**, *67*, 411–424. [[CrossRef](#)]
72. Chang, H.H.; Lu, Y.Y.; Lin, S.C. An elaboration likelihood model of consumer respond action to facebook second-hand marketplace: Impulsiveness as a moderator. *Inf. Manag.* **2020**, *57*, 103171. [[CrossRef](#)]
73. Wongkitrungrueng, A.; Assarut, N. The role of live streaming in building consumer trust and engagement with social commerce sellers. *J. Bus. Res.* **2020**, *117*, 543–556. [[CrossRef](#)]
74. Longoni, C.; Bonezzi, A.; Morewedge, C.K. Resistance to medical artificial intelligence. *J. Consum. Res.* **2019**, *46*, 629–650. [[CrossRef](#)]
75. Bagozzi, R.P.; Yi, Y.; Phillips, L.W. Assessing construct validity in organizational research. *Adm. Sci. Q.* **1991**, *36*, 421–458. [[CrossRef](#)]
76. Nunnally, J.C.; Bernstein, I.H. *Psychometric Theory*; McGraw-Hill: New York, NY, USA, 1994.

77. Fornell, C.; Larcker, D.F. Evaluating structural equation models with unobservable variables and measurement error. *J. Mark. Res.* **1981**, *18*, 39–50. [[CrossRef](#)]
78. Hayes, A.F. *Introduction to Mediation, Moderation, and Conditional Process Analysis: A Regression-Based Approach*; Guilford Publications: New York, NY, USA, 2017.
79. Sheehan, B.; Jin, H.S.; Gottlieb, U. Customer service chatbots: Anthropomorphism and adoption. *J. Bus. Res.* **2020**, *115*, 14–24. [[CrossRef](#)]
80. Choi, S.; Yi, Y.; Zhao, X. The human touch vs. AI efficiency: How perceived status, effort, and loyalty shape consumer satisfaction with preferential treatment. *J. Retail. Consum. Serv.* **2024**, *81*, 103969. [[CrossRef](#)]
81. Markovitch, D.G.; Stough, R.A.; Huang, D. Consumer reactions to chatbot versus human service: An investigation in the role of outcome valence and perceived empathy. *J. Retail. Consum. Serv.* **2024**, *79*, 103847. [[CrossRef](#)]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.