



Article

Consumer Reactions to Virtual Influencer Transgressions: How Anime-Looking and AI-Driven Influencers Are Less Vulnerable

Wei Song ^{1,†}, Siyuan Wei ^{2,†}, Zinuo Li ³, Shengliang Deng ⁴  and Yuqi Du ^{5,*}

¹ Yatai School of Business Management, Jilin University of Finance and Economics, Changchun 130117, China; songweijlu@163.com

² School of Philosophy and Sociology, Jilin University, Changchun 130012, China; weisiyuan202102@163.com

³ Curtin Business School, Curtin University Singapore, Singapore 117684, Singapore; 22495921@student.curtin.edu.au

⁴ Goodman School of Business, Brock University, St. Catharines, ON L2S 3A1, Canada; sdeng@brocku.ca

⁵ Graduate School, Guizhou University of Finance and Economics, Guiyang 550025, China

* Correspondence: dyq@mail.gufe.edu.cn

† These authors contributed equally to this work.

Abstract

Virtual influencers in diverse appearances emerged and gained popularity on virtual platforms. However, how the appearances of virtual influencers affect consumers' attitudes and reactions remained largely unexplored. Through three experimental studies, this paper examines the psychological mechanism and boundary conditions for consumer reactions to virtual influencer transgressions. The results show that consumers are less forgiving and more negative in their reactions to transgressions conducted by human-like virtual influencers compared to anime-like ones, regardless of the type of transgression or the gender of the virtual influencer (Studies 1 and 2). Additionally, the driving mechanism of the virtual influencers has a moderating effect. When consumers are informed that the virtual influencer transgression is driven by a real person rather than AI, the impact of appearance on the reactions to transgressions is aggravated (Study 3). The result shows that appearance and driving mechanism both influence consumer perceptions of the virtual influencers' agency, thereby determining the degree of reaction to transgressions.

Keywords: virtual influencer; anime-like appearance; human-like appearance; transgression; AI; mind perception; perceived agency



Academic Editors: Chenglu Wang,
Zhen Li, Xiaoling Li, Hongfei Liu and
Morgan Yang

Received: 2 June 2026

Revised: 6 July 2026

Accepted: 6 July 2026

Published: 9 July 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.

1. Introduction

In today's society, the binary interaction between businesses and consumers is gradually strengthening and digital content marketing is also on the rise [1]. Within the metaverse, virtual influencers, among other virtual beings and virtual idols, have become key actors who drive the traffic on digital platforms. Examples for virtual influencers include Lil Miquela from the United States, Shudu from the United Kingdom, Imma and Hatsune Miku from Japan, all heralding the arrival of the celebrity 2.0 era [2]. Compared to human influencers, virtual influencers exhibit greater adaptability and a wider stretch of appearance [3]. More notably, virtual influencers distinguish from traditional human influencers in terms of attractiveness, credibility, and endorsement effectiveness [4–6]. By the same token, virtual influencers' misbehavior would be more disciplined compared to human influencers [7,8], which would presumably lead to fewer transgressions and less

financial loss [9]. Consequently, explosive growth of virtual influencers has been seen as a growing number of them are employed as brand ambassadors for companies [10]. Nonetheless, virtual influencers still exhibit transgressions in certain situations [11,12], such as striking examples of identity deception by Lil Miquela and inappropriate racist remarks by Blawko [13,14]. Despite the advantages of virtual influencers over human influences in this aspect [15,16], attentions on the potential transgressions conducted by different virtual influencers have been raised [10].

Virtual influencers are categorized into human-like virtual influencers and anime-like virtual influencers based on the degree of anthropomorphism [17]. Recent studies have found differences between the two types when studying their user engagement [18], social interaction [19], brand attitude [20], and information dissemination effects [21,22]. In view of the aforementioned research gap, this paper aims to investigate the patterns of consumer reactions to transgressions conducted by the different types of virtual influencers. By examining the boundary conditions and underlying psychological mechanisms, we attempt to make two primary contributions. First, this paper takes a comprehensive look into virtual influencers based on their outside appearances and inside driving mechanisms, and reveals the interactive effects of these dominant factors on consumer reactions to transgressions. Second, based on mind perception theory, this paper explains the underlying psychological mechanisms for the differed levels of reactions to transgressions conducted by the different types of virtual influencers.

2. Theoretical Background and Hypotheses Development

2.1. Human-like Versus Anime-like Virtual Influencers

As advanced substitutes for human influencers [23], virtual influencers are digital figures running by computer imaging technology combined with machine learning, possessing significant social media influence [18]. Mouritzen et al. [9] pointed out the critical role of anthropomorphism in virtual influencers' interactions with audiences. Some virtual influencers have realistic human-like appearances, such as Lil Miquela who has 2.6 million followers on Instagram and YouTube (see example on the left in Figure 1). Others possess cartoon-like features with animation characteristics (e.g., large eyes and small mouths), such as the influential virtual singer Hatsune Miku (see example on the right in Figure 1). Recent studies accepted the human-like versus anime-like virtual influencer taxonomy [17,20].

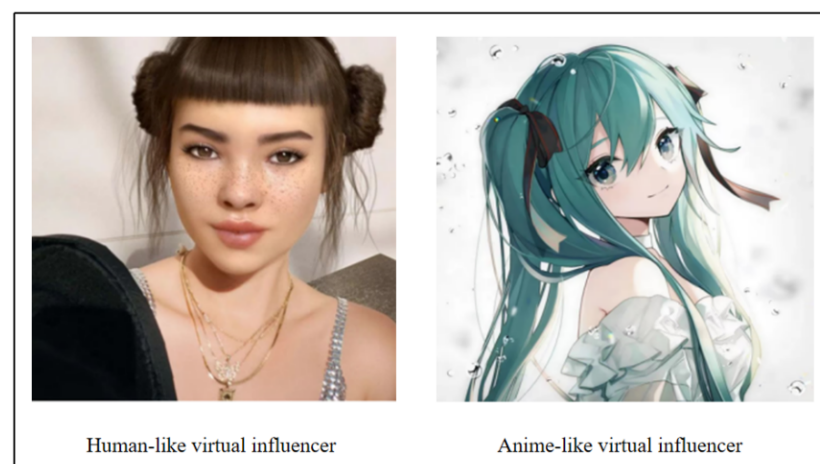


Figure 1. Taxonomy of Virtual Influencer Appearance.

According to Uncanny Valley Theory, overly realistic human-like figures evoke fear and a sense of oddness, negatively affecting consumers' emotional attachment and willing-

ness to engage [24,25]. Qu and Baek [26] found that perceived credibility of human-like figures is significantly lower than that of humans and cartoon characters. Additionally, Arsenyan and Mirowska [18] used text analysis to compare social interactions between audiences and human-like, human, and cartoon-like virtual influencers, and found that social media audience exhibits more anxiety, anger, and sadness in response to human-like virtual influencers, along with more profanity in comments. Additionally, research has shown that when consumers immersed in angry emotions, a higher degree of anthropomorphism in AI service avatars can lead to a decrease in customer satisfaction and overall evaluation of the company [27].

On the other hand, Social Response Theory suggests that people treat computers as social actors [28]. When anthropomorphic features, such as human voice and appearance, appear in non-human entities, responses tend to follow social interaction norms and exhibit positive emotions and behavior [29,30]. Xie-Carson et al. [31] found that content related to human-like figures is rated more favorably than 3D and 2D cartoon characters. More often than not, consumers rate human-like figures better in attractiveness, professionalism, and credibility, leading to stronger persuasive effects and more positive brand influence [17,20,32].

Thus, previous literature tinted with inconsistent conclusions implies both human-like and anime-like virtual influencers have certain advantages over each other, depending on the situation and research context. This study explores the effects of virtual influencer transgressions, with a focus on if human-like or anime-like appearances differ on consumer attitudes and reactions.

2.2. Virtual Influencer Transgressions and User Reactions

Transgressions conducted by brand and endorsers cause severe negative impacts, affecting both companies and endorsers [33–35]. Reinikainen et al. [36] found that influencers and brands share a symbiotic relationship where transgressions by one party can have negative spillover effects on the other, leading to lower brand attitudes, reduced trust, and negative purchase intentions. Furthermore, just the act of swearing can negatively impact both the influencer and the brand [37]. While virtual influencers appear to be less prone to transgressions due to their programmable nature, transgressions and other dreadful mistakes can still occur during interactions with audiences. The phenomena were evidenced by incidents like offensive remarks by Hololive's virtual idols Akai Haato and Kiryu Coco, which led to all Hololive virtual figures being withdrawn from the Chinese market [38].

The negative reactions caused by transgressions often lead to unfollowing the influencer, reduced purchase intentions, and even calls for punishment [39–41]. The observed negative evaluations in response to transgressions, are accompanied by forgiveness at varying levels [42,43]. Forgiveness propensity is a common measure of attitude towards service failures, indicating the degree to which users understand and forgive a virtual influencer's mistakes or transgressions [44,45]. This study will examine both forgiveness propensity and behavioral responses as dependent variables for consumer reactions to virtual influencer transgressions.

2.3. Virtual Influencer Appearance and Mind Perception

Consumer attitudes and behavioral tendencies towards virtual influencers are often explained by Mind Perception Theory [46,47]. Mind perception, also known as humanizing or mentalizing, is a cognitive process that helps individuals infer others' motives, beliefs, and intentions [48,49]. According to Mind Perception Theory [50], to determine another person's intentions or feelings, individuals seek to understand that person's mental presence, which includes two distinct aspects—agency and experience. Agency reflects the

person's ability to regulate intentions and behaviors, such as self-control and judgment and so on, while experience indicates the person's ability to understand emotions and feelings, such as hunger, fear, or pain. The mind perception route is also applied when people attempt to understand human-comparable minds, such as groups, animals, and mechanical agents, with a motivation to establish social connections [51]. Therefore, even though lifeless objects like virtual influencers lack thoughts, consumers tend to habitually attribute human-like mental capacity to them.

Prior studies have found that consumers tend to perceive lower mental capacity in virtual influencers than human influencers [16,52]. This is because consumers believe that algorithm-driven non-human entities lack empathy and cognitive abilities [46,53]. When an interaction with non-human takes place, the degree of anthropomorphism is a crucial indicator of mental capacity [54]. For instance, human characteristics and human behaviors of robot avatars are associated with stronger mind perception judgments [55,56]. On virtual platforms, facial characteristics of human appearance provide huge cues for consumers' mind perception of virtual beings [57]. Thus, consumers' stronger mind perception can be found in human-like virtual influencers due to the higher levels of anthropomorphism, especially based on facial features.

Consumers' mind perception of virtual influencers shapes their attitudes and behavioral reactions. Ham et al. [58] conducted a content analysis of avatar and human interactions on Instagram, and found that emotional texts presented by avatars can enhance audience evaluations on their emotional intelligence, leading to the audience's positive attitudes and improved interactions. Additionally, Yang et al. [59] also indicates that an increase in the behavioral anthropomorphism of virtual influencer will enhance consumers' evaluation of their capabilities and lead to more positive consumption behaviors. In contrast, lower mind perception in virtual influencers leads to a series of negative reactions are seen, such as lower willingness to share, decreased brand attitudes, and weaker purchase intentions [46,52]. When dealing with transgressions conducted by the different appearance types of virtual influencers, does stronger mind perception in human-like virtual influencers lead to better or worse consumer reactions to transgressions? Gray et al. [60] noted that objects associated with higher mental capacity are attributed with more moral rights and responsibilities. This is because those who have cognitive autonomy are believed to make their own decisions and take conscious actions, and thus are held more accountable for their mistakes. This notion is applicable to virtual influencers [61]. When certain virtual influencers are perceived to lack cognitive abilities, the moral demand for responsibility decreases and forgiveness propensity increases [62]. Since human-like virtual influencers are perceived to have greater control over their behaviors than anime-like ones, consumers tend to give less forgiveness towards their transgressions. Stronger mind perception eventually results in lower forgiveness propensity and more negative behavioral responses. Therefore, we propose the following hypothesis:

H1: *Compared to anime-like virtual influencers, individuals exhibit lower forgiveness propensity and more negative behavioral responses to transgressions by human-like virtual influencers.*

2.4. Driving Mechanism as Boundary Condition

Human appearance is not the only cue consumers can rely on with regard to mind perception. Park and Sung [17] explained that virtual influencers are operated on different technologies. Some virtual influencers are entirely driven by digital imaging technologies and algorithms, such as Lil Miquela and the cartoon character Apoki. Others are merely virtual appearances attached to a real person, using facial expressions, voice, and motion capture technologies to project a real person's behavior into a virtual figure, such as Japanese virtual influencer Kizuna AI and the virtual celebrity Monica Quin on social

networking app Zepeto. In general, the driving mechanisms for virtual influencers are either AI-driven or real person-driven [17,63]. Real person-driven was akin to avatar role-playing. This kind of virtual influencers were created through real people wearing body and facial motion capturing equipment combined with voice synthesis technology [64]. Alongside the advancement of digital technologies, AI-driven virtual influencers have also become mainstream. The operation of AI-driven virtual influencers relies on high-precision multimodal technology and deep learning for movement restoration, akin to a ChatGPT 5.4 Pro (OpenAI, San Francisco, CA, USA) with a physical image [65,66].

Prior studies have not reached a conclusion on whether digital humans driven by real person are better or worse than those driven by AI. On the one hand, some scholars suggested that real person-driven virtual influencers yield better interaction outcomes than AI-driven ones. For example, consumers experience enhanced flow and more positive emotions when interacting with human avatars compared to computers [67]. Consumers think of real person-driven virtual influencers as having stronger social presence on social media platforms, which shows stronger persuasive effectiveness than do computerized figures [68]. On the other hand, a study by Seymour et al. [69] found that the type of driving mechanism does not affect consumers' ratings of affinity and trustworthiness towards virtual influencers.

On the flip side, when consumers encounter virtual influencer transgressions, they might perceive real person-driven virtual influencers under stronger mind perception. Although there is no relevant evidence in the field of virtual influencers, research on service robots suggests that consumers are more likely to express anger towards human agents (vs. AI agents) in the context of service failures and attribute more responsibility to the customer service itself [70]. The study by Garvey et al. [71] also indicates that, compared with human agents, consumers express less dissatisfaction with bad news delivered by AI agents. Since both appearance and driving mechanism are commonly used to judge the mental capacity of virtual influencers, we expect an interaction effect by the two. When consumers are provided with both appearance and driving mechanism of virtual influencers, they will exhibit a greater degree of negative reactions to the transgressions conducted by human-like virtual influencers who are real person-driven. In contrast, when informed that the anime-like virtual influencers are AI-driven, consumers perceive the lowest level of mental capacity in them, thus leading to highest level of forgiveness propensity and least negative behavioral responses. According to Mind Perception Theory, we propose the following hypotheses:

H2: *Driving mechanism moderates the effects of influencer appearance on consumer reactions to virtual influencer transgressions.*

H3: *Mind perception mediates the interactive effect of influencer appearance and driving mechanism on consumer reactions to virtual influencer transgressions.*

This study's theoretical framework is shown in Figure 2. The rest of the paper is organized as follows: Study 1 uses existing human-like and anime-like virtual influencer images in the experiment to explore the main effects of consumer reactions to virtual influencer transgressions (H1), taken into consideration the potential impact of transgression types. Study 2 uses self-designed virtual influencer images to verify the robustness of H1, taken into consideration the potential impact of virtual influencer gender. Finally, Study 3 examines the boundary conditions of driving mechanisms (H2) and the mediating effects of mind perception (H3).

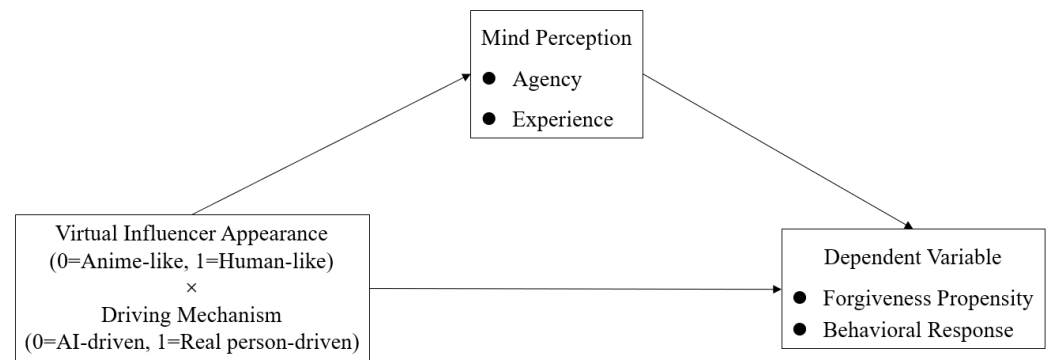


Figure 2. Theoretical Framework.

3. Study One: Consumer Reactions to Transgressions Based on Virtual Influencer Appearance

The purpose of Study 1 is to explore the main effects of virtual influencer appearance on consumer reactions to their transgressions. Virtual influencer transgressions are often observed in many ways, such as inappropriate product endorsements and offensive remarks during live broadcasts. We choose two distinct kinds of transgressions as experimental scenarios—false endorsement and discrediting statements.

3.1. Participants and Design

Study 1 employed a 2 (influencer appearance: human-like vs. anime-like) $\times$ 2 (transgression behavior type: false endorsement vs. discrediting statements) between-subjects design and Credamo (Beijing, China; Chinese equivalents of Mturk for online data gathering) was used to recruited participants. Finally, 200 ($M_{\text{age}} = 29.96$, $SD = 6.64$, 35% males) subjects with no outliers, normal reaction times and who pass the attention test were recruited at random and then were distributed to one of the four conditions. Table 1 reports the participants' demographic profile. By using G*Power 3.1 (Heinrich Heine University, Düsseldorf, Germany) to analyze the statistical power of the sample size [72], when the $f = 0.5$, $df = 1$, $\alpha = 0.05$ with $N = 200$, the Power $(1 - \beta) > 0.99$. This indicates the sample size for this experiment has good statistical power.

Table 1. Participants Information in Study 1 ($N = 200$).

Demographic Variable	Category	Number (Proportion)
Age	18–24	40 (20%)
	25–34	124 (62%)
	35–44	27 (13.5%)
	45–54	9 (4.5%)
Gender	Male	70 (35%)
	Female	130 (65%)
Education Background	High school or lower	17 (8.5%)
	Some college	31 (15.5%)
	College graduate	92 (46%)
	Postgraduate	60 (30%)

3.2. Procedure

In Study 1, we used existing images to represent human-like and anime-like influencers. As shown in Figure 3, the image of virtual influencer “Higaga” (created by Baidu Corporation, Beijing, China) was used as the human-like virtual influencer and the image of another virtual influencer “Du Xiaoxiao” (also created by Baidu Corporation) was used

as the anime-like virtual influencer. We conducted a pretest ($N = 60$; $M_{\text{age}} = 24.78$, $SD = 4.88$, 38.30% males) to measure the attractiveness of the images (1 = not attractive at all, 7 = very attractive) and found that the two influencers did not significantly differ in attractiveness ($t(58) = 0.025$, $p = 0.922$).

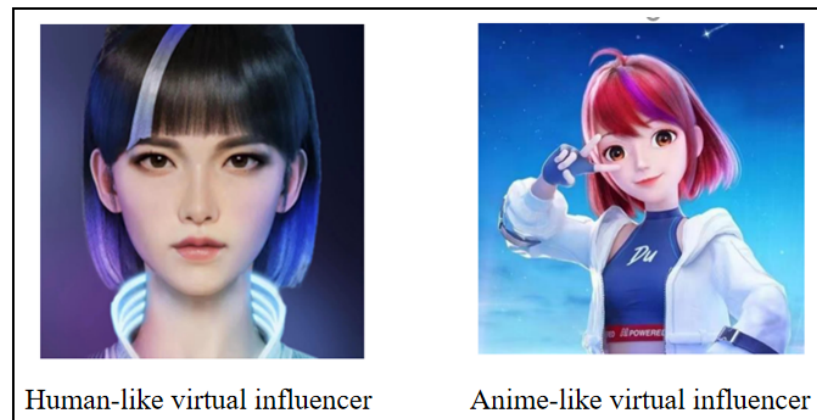


Figure 3. Experimental Stimuli in Study 1.

In the cover story, we presented a brief explanation of the influencers, emphasizing their human-like or anime-like appearance to enhance the manipulation (see Appendix A for details). Next, the participants were informed of one of the two transgression behaviors. In the false endorsement scenario, the participants were described a scenario where a virtual influencer recommended a product to the fans, but the fans later discovered that the product was of poor quality. In the discrediting statements scenario, the participants were described a situation in which a virtual influencer made discriminatory remarks about disabled individuals (see Appendix B for details).

After that, the participants were asked to imagine themselves as fans of the influencer and indicate their evaluations of the influencer. We used forgiveness propensity and behavioral response as dependent variables. Forgiveness propensity was tested by three 7-point Likert items adapted from Zhao et al. [16] ($\alpha = 0.91$): (1) how likely are you to forgive the influencer (1 = very unlikely, 7 = very likely); (2) I think I can forgive the mistake made by the influencer (1 = strongly disagree, 7 = strongly agree); and (3) if the influencer can fix the transgression in time, I am willing to forgive them (1 = strongly disagree, 7 = strongly agree). Behavioral response was measured using three 7-point Likert items adapted from Zhao et al. [16] ($\alpha = 0.79$; 1 = strongly disagree, 7 = strongly agree): (1) I have a bad impression of the influencer; (2) I will continue to follow the influencer (reversed item); and (3) I would not buy any products recommended by the influencer anymore.

3.3. Results

2 (influencer appearance) $\times$ 2 (transgression behavior type) two-way ANOVAs on forgiveness propensity and behavioral response showed significant main effects of influencer appearance ($F_s > 13.80$, $p_s < 0.001$) and transgression behavior type ($F_s > 6.99$, $p_s < 0.01$). Meanwhile, the interactions were not significant ($F_s < 0.01$, $p_s > 0.94$), indicating that the effects of influencer appearance on forgiveness propensity and behavioral response are consistent across different transgression behavior types (see Figure 4). As predicted, the anime-like virtual influencer is associated with higher forgiveness ($M_{\text{anime-like}} = 4.17 \pm 1.65$, $M_{\text{human-like}} = 3.44 \pm 1.40$, $t(198) = -3.37$, $p < 0.01$, $d = 0.48$) and less negative behavioral response ($M_{\text{anime-like}} = 4.04 \pm 1.68$, $M_{\text{human-like}} = 4.97 \pm 1.31$, $t(198) = 4.33$, $p < 0.001$, $d = 0.62$) than the human-like virtual influencer. H1 was supported.

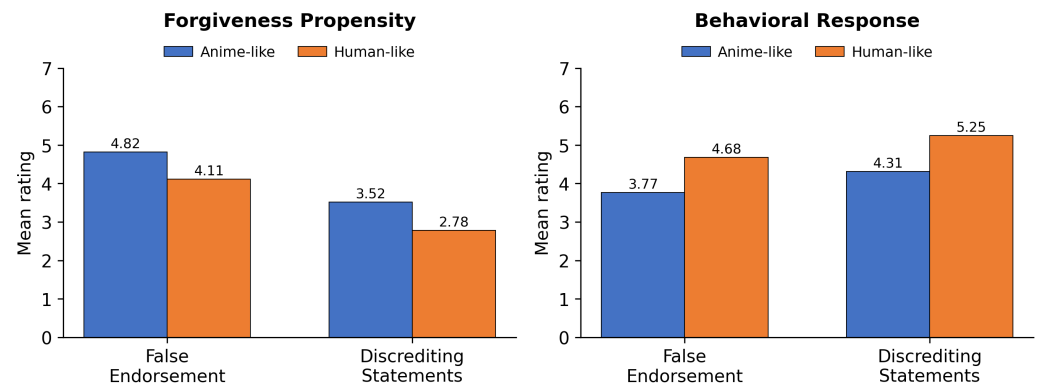


Figure 4. Results of Study 1.

3.4. Discussion

The results of Study 1 provided initial evidence that the appearance of virtual influencers matters when consumers encounter virtual influencer transgressions. Particularly, it is more probable that anime-like virtual influencers will receive forgiveness from consumers, regardless of the type of transgression. Study 1 rules out the potential impact of transgression types.

4. Study Two: Main Effects Based on Virtual Influencer Gender

The results in Study 1 revealed the main effects of virtual influencer appearance on consumer reactions to their transgressions, and demonstrated the consistency of the main effects across different types of transgressions. However, the experimental material used in Study 1 were both female. Although female virtual influencers are prevalent on virtual platforms, there are also a good number of male virtual influencers. Given that the appearances between male and female influencers differ considerably, the gender of the influencer might be a potential confounding variable. Therefore, Study 2 employed a set of human-like and anime-like images for both female and male virtual influencers to examine the consistency of the main effects across virtual influencer genders. To control potential confounding effects due to existing image characteristics, we created fictitious images as experimental stimuli in Study 2.

4.1. Participants and Design

Study 2 adopted a 2 (influencer appearance: human-like vs. anime-like) $\times$ 2 (influencer gender: male vs. female) between-subjects design. A total of 200 participants ($M_{\text{age}} = 30.18$, $SD = 7.36$, 44.0% males) were recruited via Credamo like study 1 and randomly assigned them to one of the four conditions. Table 2 reports the participants' demographic profile.

4.2. Procedure

Following the same procedure used in Study 1, the participants were presented with a brief description of the influencer paired with a fictitious image (see Figure 5). The participants were then informed of the false endorsement transgression behavior used in Study 1. After that, we tested the participants' forgiveness propensity ($\alpha = 0.86$) and behavioral responses ($\alpha = 0.87$) as measured in Study 1.

4.3. Results

2 (influencer appearance) $\times$ 2 (influencer gender) two-way ANOVAs on forgiveness propensity and behavioral response showed neither significant main effects of influencer gender nor significant interactions ($F_s < 3.40$, $p_s > 0.067$). The main effects of influencer appearance were significant ($F_s > 6.56$, $p_s < 0.011$, see Figure 6). Further comparisons showed that participants in the anime-like virtual influencer condition indicated

higher forgiveness ($M_{\text{anime-like}} = 4.87 \pm 1.16$, $M_{\text{human-like}} = 4.41 \pm 1.39$, $t(198) = -2.56$, $p = 0.011$, $d = 0.36$) and less negative behavioral response ($M_{\text{anime-like}} = 3.64 \pm 1.35$, $M_{\text{human-like}} = 4.25 \pm 1.49$, $t(198) = 3.06$, $p = 0.003$, $d = 0.43$) than human-like virtual influencers. H1 was supported again.

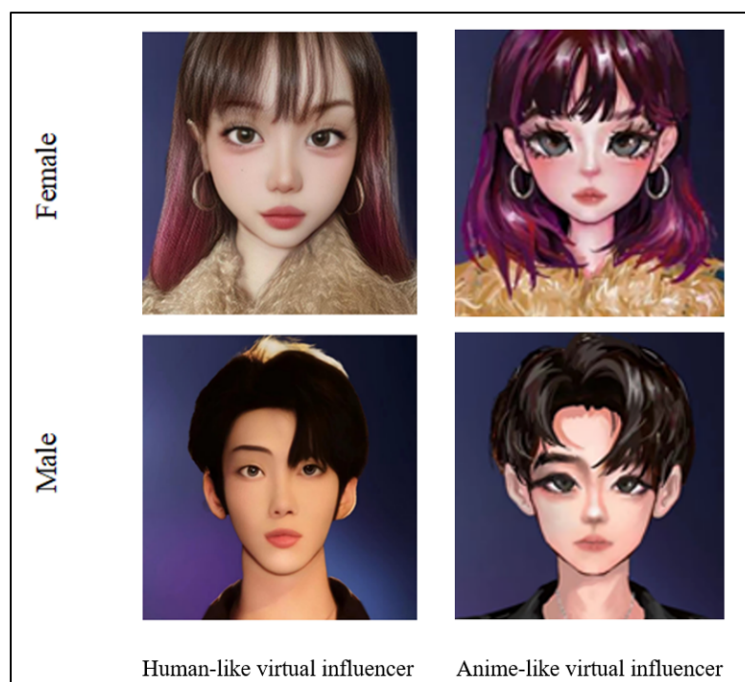


Figure 5. Experimental Stimuli in Study 2.

Table 2. Participants Information in Study 2 ($N = 200$).

Demographic Variable	Category	Number (Proportion)
Age	18–24	47 (23.5%)
	25–34	106 (53%)
	35–44	37 (18.5%)
	45–54	7 (3.5%)
	55–64	2 (1%)
Gender	Male	88 (44%)
	Female	112 (56%)
Education Background	High school or lower	11 (5.5%)
	Some college	37 (18.5%)
	College graduate	94 (47%)
	Postgraduate	58 (29%)

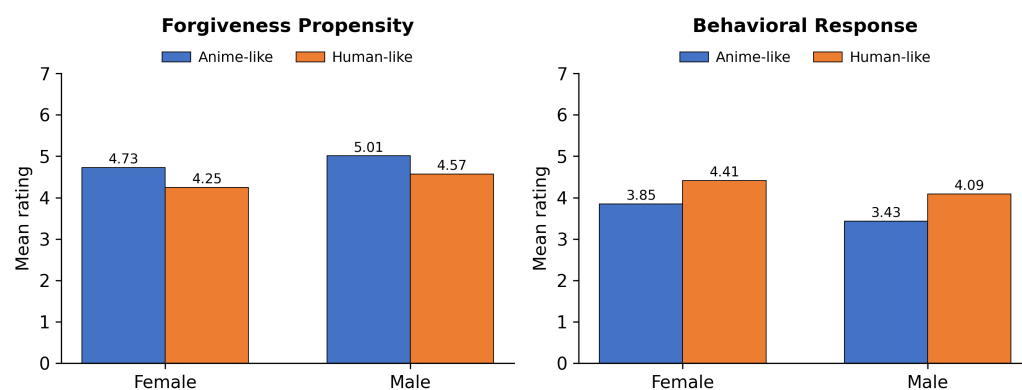


Figure 6. Results in Study 2.

4.4. Discussion

Results in Study 2 confirmed the main effects hypothesized in H1 and clarified that there is no effect of influencer gender. Regardless of the gender of the virtual influencer, their appearance significantly impacts consumers' attitudes and behaviors following transgressions. Consumers exhibit lower forgiveness propensity and more negative behavioral responses to transgressions conducted by human-like virtual influencers compared to anime-like virtual influencers. Therefore, Study 3 uses only female virtual influencers as experimental stimuli.

5. Study Three: Moderating Effects of Driving Mechanism and Mediating Effects of Mind Perception

The primary purpose of Study 3 was to examine the boundary conditions of the driving mechanism of virtual influencers as well as the mediating effects of mind perception. As mind perception includes two facts, Study 3 measured both agency and experience to see how they explain consumer reactions to virtual influencer transgressions.

5.1. Participants and Design

Study 3 adopted a 2 (influencer appearance: human-like vs. anime-like) $\times$ 2 (driving type: AI-driven vs. real person-driven) between-subjects design. We also recruited 200 participants from Prolific (Prolific, London, UK) at random. Owing to the problems like response time and attention check, study 3 included 195 valid subjects ($M_{\text{age}} = 36.96$, $SD = 13.14$, 45.6% males), and they also were given one of the four conditions in a random manner. Table 3 reports the participants' demographic profile.

Table 3. Participants Information in Study 3 ($N = 195$).

Demographic Variable	Category	Number (Proportion)
Age	18–24	34 (17.4%)
	25–34	67 (34.4%)
	35–44	41 (21%)
	45–54	34 (17.4%)
	55–64	11 (5.6%)
	>65	8 (4.1%)
Gender	Male	103 (52.8%)
	Female	89 (45.6%)
	Non-binary/Third gender	3 (1.5%)
Education Background	High school or lower	18 (9.2%)
	Some college	44 (22.6%)
	College graduate	88 (45.1%)
	Postgraduate	45 (23.1%)
Region	Australia/Oceania	71 (36.4%)
	Africa	10 (5.1%)
	North America	73 (37.4%)
	Europe	39 (20%)
	Asia	2 (1%)
Social Media Usage Frequency	Less than once a week	5 (2.6%)
	Once a week	2 (1.0%)
	A few times a week	16 (8.2%)
	Once a day	28 (14.4%)
	Several times a day	144 (73.8%)

5.2. Procedure

First, the participants were presented a brief overview of the influencer, matched with a picture (see Appendix C for details). We chose a human-like virtual influencer image and an anime-like virtual influencer image from Bilibili, a popular video platform in China (see Figure 7). To mitigate potential confounding effects, we standardized the number of followers and information descriptions of the virtual influencers.

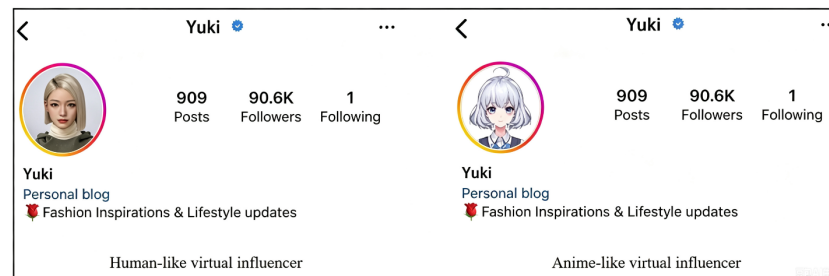


Figure 7. Experimental Stimuli in Study 3.

We manipulated the driving mechanisms of the virtual influencer by providing an AI-driven scenario and a real person-driven scenario to the participants. The participants were then given a transgression behavior scenario involving discrediting statements (see Appendix B for details). After that, the participants answered an attention check question. Those who passed the attention check were asked to indicate their forgiveness propensity ($\alpha = 0.94$) and behavioral response ($\alpha = 0.83$) using the same measures in Studies 1 and 2. Additionally, the participants reported their mind perception on the influencer. Mind perception was tested by a 7-point Likert scale on twelve items (“To what extent do you think the influencer can...”; 1 = not at all, 7 = extremely; $\alpha = 0.95$; [73]), including six agency-related capacities (“communicate with others,” “is capable of thinking,” and so on; $\alpha = 0.92$) and six experience-related capacities (“sensitive to pain,” “experience happiness,” “experience fear,” and so on; $\alpha = 0.98$).

5.3. Results

Interaction Effects on Consumer Reactions to Transgressions. 2 (influencer appearance) $\times$ 2 (driving mechanism) two-way ANOVAs on forgiveness propensity and behavioral response showed no significant main effects of influencer appearance or driving mechanism ($F_s < 3.61$, $p_s > 0.059$). Still, the interactions were significant ($F_s > 4.82$, $p_s < 0.029$). H2 was supported.

For dependent variable forgiveness propensity, the interaction of influencer appearance and driving mechanism was significant ($F(1, 191) = 5.35$, $p = 0.022$, $\eta_p^2 = 0.027$). Further analysis of simple effects showed that the difference between human-like ($M = 2.99$, $SD = 1.49$) and anime-like ($M = 2.85$, $SD = 1.18$) for AI-driven virtual influencers was not significant ($F(1, 191) = 0.21$, $p = 0.651$, $\eta_p^2 = 0.001$). However, consumers were more inclined to forgive the anime-like virtual influencer ($M = 3.01$, $SD = 1.95$) than the human-like virtual influencer ($M = 2.16$, $SD = 1.07$) when encountering real person-driven virtual influencers ($F(1, 191) = 7.92$, $p = 0.005$, $\eta_p^2 = 0.040$, see Figure 8).

For dependent variable behavioral response, the interaction of influencer appearance and driving mechanism was significant ($F(1, 191) = 4.83$, $p = 0.029$, $\eta_p^2 = 0.025$). Further analysis of simple effects showed that the difference between human-like ($M = 5.45$, $SD = 1.50$) and anime-like ($M = 5.51$, $SD = 1.36$) AI-driven virtual influencers was not significant ($F(1, 191) = 0.04$, $p = 0.834$, $\eta_p^2 = 0.000$). However, consumers react more negatively to the human-like virtual influencer ($M = 6.18$, $SD = 0.68$) than the anime-like virtual influencer

($M = 5.38$, $SD = 1.67$), when the virtual influencers are real person-driven ($F(1, 191) = 8.37$, $p = 0.004$, $\eta_p^2 = 0.042$, see Figure 9).

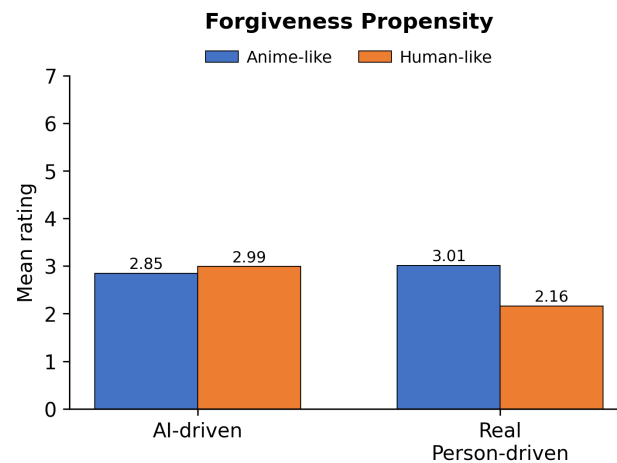


Figure 8. Results of Interaction Effect on Forgiveness Propensity in Study 3.

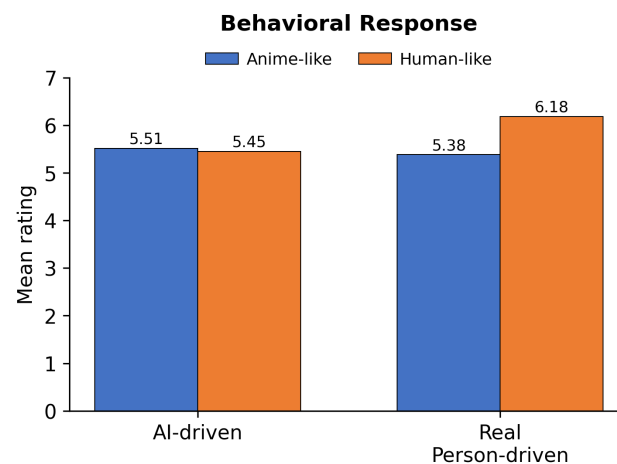


Figure 9. Results of Interaction Effect on Behavioral Response in Study 3.

The results suggest that the main effects found in Studies 1 and 2 do not exist when consumers know the virtual influencers are operated by AI. To investigate the mediating effects of mind perception, we further conducted an analysis with the two subdimensions of mind perception as dependent variables.

Interaction Effects on Mind Perception (Agency and Experience). We conducted a 2 (influencer appearance) $\times$ 2 (driving mechanism) two-way ANOVA on the two dimensions of mind perception, agency and experience, respectively. The results with agency as the dependent variable showed that the main effect of driving mechanism was significant ($F(1, 191) = 146.21$, $p < 0.001$, $\eta_p^2 = 0.434$), indicating that consumers consider AI-driven virtual influencers a lower level of agency compared to real person-driven ones. The main effect of influencer appearance was not significant ($F(1, 191) = 0.72$, $p = 0.397$, $\eta_p^2 = 0.004$), but the interaction was significant ($F(1, 191) = 4.79$, $p = 0.030$, $\eta_p^2 = 0.024$). Further analysis of simple effects showed that the difference between human-like ($M = 2.70$, $SD = 1.14$) and anime-like ($M = 2.92$, $SD = 1.26$) virtual influencers was not significant when they are AI-driven ($F(1, 191) = 0.90$, $p = 0.344$, $\eta_p^2 = 0.005$). However, consumers perceive a higher level of agency in human-like ($M = 4.95$, $SD = 0.89$) than anime-like virtual influencers ($M = 4.47$, $SD = 1.06$) when they are real person-driven ($F(1, 191) = 4.60$, $p = 0.033$, $\eta_p^2 = 0.024$, see Figure 10). This interactive effect is consistent with the patterns found when the forgiveness propensity and behavioral response are taken as the dependent variables.

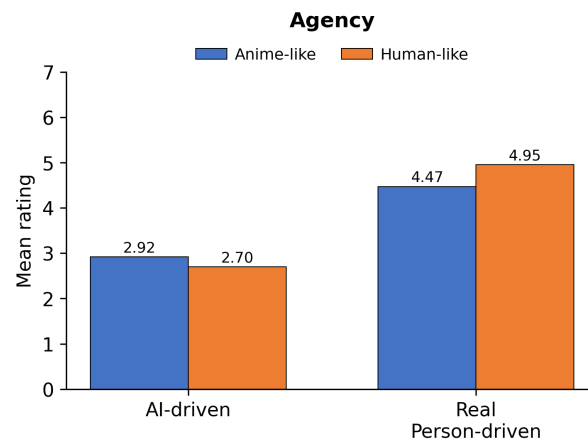


Figure 10. Results of Interaction Effect on Agency in Study 3.

When experience was used as dependent variable, only the main effect of driving mechanism was significant ($F(1, 191) = 152.90, p < 0.001, \eta_p^2 = 0.445$). The main effect of influencer appearance was not significant ($F(1, 191) = 1.16, p = 0.283, \eta_p^2 = 0.006$), and the interaction effect was also not significant ($F(1, 191) = 1.34, p = 0.249, \eta_p^2 = 0.007$, see Figure 11). This suggests that agency, and not experience, is the underlying mechanism explaining consumer reactions to transgressions conducted by virtual influencers. To verify this hypothesis, we conducted a mediated moderation analysis.

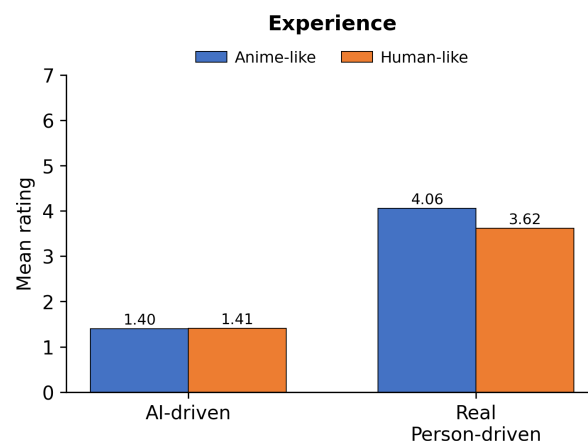


Figure 11. Results of Interaction Effect on Experience in Study 3.

Moderated Mediation Effects. Influencer appearance (0 = anime-like virtual influencer, 1 = human-like virtual influencer) and driving mechanism (0 = AI-driven, 1 = real person-driven) were coded as dummy variables. Using Model 8 of the PROCESS macro in SPSS Statistics 27 (IBM Corp., Armonk, NY, USA) [74], we ran moderated mediation analysis with agency and experience as parallel mediators and forgiveness propensity and behavioral response as the outcome variables.

The results are shown in Table 4. The path analysis with forgiveness propensity as the dependent variable reveals that the interaction between influencer appearance and driving mechanism positively affects agency ($\beta = 0.687, p = 0.030$), and agency negatively affects forgiveness propensity ($\beta = -0.473, p < 0.001$). However, when experience was used as a mediator, only the effect of experience was significant ($\beta = 0.242, p < 0.01$). The results indicate that only the mediation effect of agency is significant (see Figure 12).

When behavioral response is used as the dependent variable, the path analysis results show that the interaction between influencer appearance and driving mechanism positively

affects agency ($\beta = 0.687, p = 0.030$), and agency positively affects behavioral response ($\beta = 0.291, p < 0.001$). However, when experience was used as a mediator, only the effect of experience was significant ($\beta = -0.202, p < 0.01$). Thus, the mediation effect of agency is significant, and the mediation effect of experience is not (see Figure 13). Overall, the moderated mediation analysis results showed not only the interaction effects between influencer appearance and driving mechanism on consumer reactions, but also that perceived agency plays the mediating role in the relationship. H3 was partially supported.

Table 4. Test of Moderated Mediation Effects.

Dependent Variables	Mediation	Moderation	Effect	BootSE	BootLLCI	BootULCI
Forgiveness Propensity	Agency	AI-driven	0.10	0.12	−0.13	0.37
		Real person-driven	−0.23	0.10	−0.45	−0.04
	Experience	AI-driven	0.003	0.04	−0.08	0.09
		Real person-driven	−0.11	0.10	−0.34	0.06
Behavioral Response	Agency	AI-driven	−0.06	0.09	−0.27	0.09
		Real person-driven	0.14	0.08	0.03	0.04
	Experience	AI-driven	−0.003	0.03	−0.07	0.06
		Real person-driven	0.09	0.08	−0.04	0.27

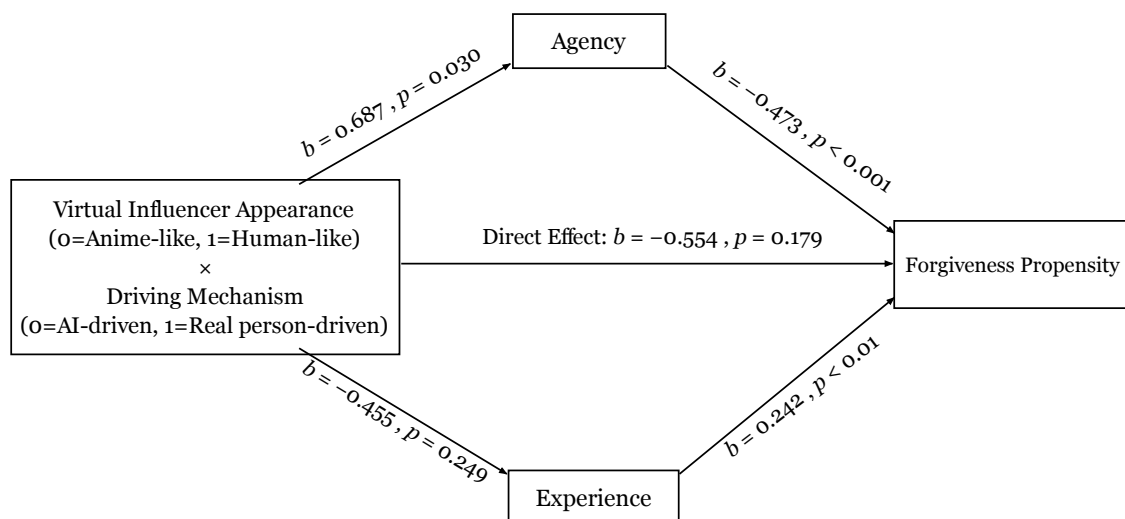


Figure 12. Mediation Model with Forgiveness Propensity as Dependent Variable.

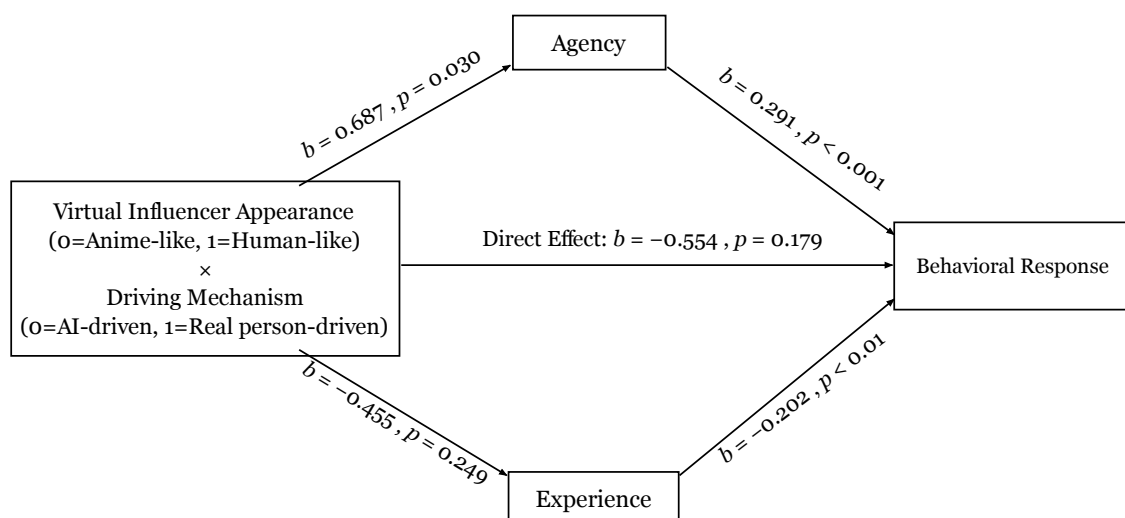


Figure 13. Mediation Model with Behavioral Response as Dependent Variable.

5.4. Discussion

Extending the findings in Studies 1 and 2, the results in Study 3 revealed that virtual influencer appearance and driving mechanism form an interaction effect on consumer reactions to transgressions. Additionally, Study 3 showed that the interaction effect is mediated by perceived agency, one dimension of mind perception, but not the other dimension, perceived experience. The findings suggest that the different dimensions of mind perception play distinct roles in consumers' information processing. These results support the research hypothesis H2 and H3. More specifically, when virtual influencers are AI-driven, the algorithmic nature perceived by consumers lead them to discount appearance information and directly infer low agency, resulting in higher forgiveness propensity and more less negative responses toward transgressions committed by AI. Conversely, real person-driven virtual influencers tend to be considered high agency due to its people-run nature. Human-like appearance and real person driving mechanism together prompt users to perceive the highest levels of human-like mental capacity behind these virtual influencers, resulting in the lowest forgiveness propensity and the most negative behavioral response.

6. General Discussion

Through three experimental studies, this paper illustrates how the appearance (human-like vs. anime-like) and driving mechanism (real person-driven vs. AI-driven) of virtual influencers jointly affect consumers' forgiveness propensity and behavioral response after transgressions take place. Consumers tend to exhibit lower forgiveness propensity and less positive behavioral response towards human-like virtual influencers compared to anime-like ones. However, when provided with driving mechanisms for the virtual influencers, consumers tend to show higher forgiveness propensity and more positive behavioral response towards AI-driven virtual influencers, regardless of their appearance. In contrast, real person-driven virtual influencers receive lower forgiveness propensity and less positive behavioral response to transgressions. Furthermore, when consumers are informed that the human-like virtual influencers are operated by a real person, they exhibit the lowest level of forgiveness propensity and the most negative behavioral response to transgressions. The findings indicate that consumers take into consideration both the external appearance and internal instrument when judging virtual influencers' motives and behavior. The patterns remained robust across different transgression scenarios and for both male and female virtual influencers.

6.1. Theoretical Implications

Our findings extend the extant literature in several aspects.

First, this study extends anthropomorphism theory by showing that its documented benefits are not unconditional. Prior research has generally treated higher anthropomorphism as advantageous, whether conveyed through appearance [17,20,24,75] or through the presence of a real human operator [67,68], associating both with greater trust, social presence, and user satisfaction. Our findings show that these benefits reverse once a transgression occurs: anime-like, lower-anthropomorphism virtual influencers receive more forgiveness and less negative behavioral response than human-like ones, and AI-driven virtual influencers elicit less negative reactions than real person-driven ones. Higher anthropomorphism and real person-driven operation, in other words, can each become a liability rather than an asset once things go wrong.

Second, this study extends Mind Perception Theory [50,76] by clarifying why each cue shapes mind attribution and how the two combine. Appearance conveys mind indirectly, through surface-level anthropomorphic cues such as facial humanness [57]. Driving mechanism conveys mind more directly, by specifying whether an actual human mind, or an

algorithm, operates behind the persona, speaking to the presence of a mind rather than merely its outward likeness; this is consistent with driving mechanism, rather than appearance, producing the dominant effect on perceived agency in Study 3. Notably, consumers appear to weight these two cues unequally rather than combining them additively. When a virtual influencer is AI-driven, this cue alone already signals a low level of mind, leaving little for appearance to add; consumers accordingly relied less on appearance, which had no discernible effect under AI-driven framing. When a virtual influencer is real person-driven, the driving-mechanism cue is less conclusive on its own, since a real person can operate either a human-like or an anime-like avatar; appearance therefore regains its diagnostic value, and its effect on perceived agency and on consumer reactions emerged only in this condition. Appearance and driving mechanism thus do not shape agency independently: a highly diagnostic cue reduces consumers' reliance on a less diagnostic one. Experience follows a different logic, taken up next.

Third, this study extends the two-dimensional structure of mind perception theory [48] by showing that agency and experience are not interchangeable predictors of post-transgression consumer reactions, but instead operate through distinct pathways. Scholars usually concentrate on the experience on AI figures, viewing it as the main distinction between algorithms and humans [52,77]. Other studies also found that consumers consider virtual influencers with poorer emotional processing capabilities [78]. However, our findings indicated that agency, rather than experience, better explains how consumers process information over virtual influencers before they respond to virtual influencer transgressions. Human-like appearance and real person-driven mechanism are perceived to have higher agency, leading to more negative attitudes in response to transgressions. Conversely, AI-driven and anime-like virtual influencers are seen as having diminished levels of agency, resulting in higher forgiveness propensity and less negative responses. Notably, both agency and experience significantly predicted these outcomes; the two dimensions differ not in whether they matter, but in whether their contribution is contingent on the appearance-by-driving-mechanism interaction. Only agency was itself significantly shaped by this interaction, whereas experience maintained a significant but unconditional relationship with both outcomes. This pattern points to consumers' cognitive reasoning about culpability, rather than their emotional response to the target, as what drives the interaction effect specifically. Specifically, agency functions as a culpability signal diagnostic of who controlled the transgressive act, which is precisely the information driving mechanism conveys, explaining why agency's effect is contingent on driving mechanism; experience, by contrast, reflects a more stable impression of the entity's capacity to suffer social consequences, eliciting a baseline sympathy response that is largely independent of who was in control. This dissociation is consistent with dyadic morality theory's separation of moral-agent and moral-patient judgments [60], and aligns with recent evidence that AI sources reduce empathy through increased psychological distance regardless of surface-level appearance cues [79]. The consistency of experience's direct path across both forgiveness propensity ($\beta = 0.242$) and behavioral response ($\beta = -0.202$) further supports a robust, generalizable sympathy-based mechanism rather than an isolated artifact.

6.2. Practical Implications

The findings offer several practical implications for how brands select, manage, and communicate about virtual influencers, particularly in circumstances involving reputational risk.

First, higher degrees of anthropomorphism for virtual influencers do not always elicit positive consumer reactions. In circumstances such as transgressions or mistakes made by pivotal virtual beings, lower anthropomorphism (for example, anime-like appearance) can mitigate the negative situation, whereas higher degrees of anthropomorphism result

in more negative consequences, consistent with the broader finding that the social signals conveyed by a virtual persona's degree of realism shape consumer evaluations [80]. This advantage should be especially valuable on platforms involving controversial topics, where transgressions and misunderstandings are more likely to occur, and it will help marketing managers deploy appropriate virtual influencer images as part of an integrated marketing communication strategy with an effective response system for potential service failures.

Second, virtual influencers' driving mechanism plays an important role in consumers' information processing when they attempt to understand the conduct of virtual influencers, and this carries a concrete implication for disclosure strategy. Because AI-driven virtual influencers are attributed lower agency and judged less harshly following a transgression, disclosing that a virtual influencer is AI-operated is likely to buffer negative reactions during a crisis, whereas disclosing that a real person operates the influencer may instead amplify backlash, consistent with broader evidence that what a brand discloses about the agent behind its communications shapes how consumers receive them [81,82]. Marketing managers should therefore treat disclosure of driving mechanism as a strategic choice to be planned before a crisis occurs, not a formality decided in the moment. The findings also remind marketing managers to more carefully select the real person who runs a real person-driven virtual influencer, as that person's misconduct is less forgivable in the eyes of consumers than an AI's.

Third, these findings also inform the ex-ante choice of which type of virtual influencer to deploy in the first place, rather than only how to respond after a transgression has occurred. Xin, Hao, and Xie [83] show that companies rationally shift toward virtual influencers as the perceived risk of endorser scandal increases; the present findings refine this logic by showing that it is not virtual-versus-human status alone that matters, but the specific combination of appearance and driving mechanism, with anime-like, AI-driven personas offering the greatest protection. Brands operating in reputation-sensitive categories, or running campaigns on topics prone to public scrutiny, may therefore wish to favor this combination proactively as a form of reputational insurance built into brand strategy from the outset [84,85]. Selecting a lower-risk virtual influencer profile from the outset may thus function as a complement to, rather than a substitute for, having an effective crisis response plan in place.

6.3. Limitations and Future Research Directions

A number of limitations can be addressed by future studies.

First, this study's operationalizations carry some methodological limitations. The experiments manipulated virtual influencers using static images and textual descriptions rather than a live, interactive format; future studies should extend these findings to dynamic virtual platforms, such as livestreaming commerce where virtual streamers increasingly operate [86], to establish external validity beyond the present stimulus format.

Second, this study's boundary conditions remain only partially understood. Although Studies 1 and 2 sampled Chinese consumers and Study 3 sampled a predominantly Western one, with the appearance effect (H1) and the appearance-by-driving-mechanism interaction (H2, H3) each replicating in their respective samples, this treats culture as a check on cross-contextual stability rather than as a potential moderating factor in its own right. A study formally crossing culture as a factor within a single design, following approaches such as Wei, Syahrivar, and Simay [87], would be needed to test this directly, particularly given prior evidence of both similarities and differences in how Chinese and Western consumers assign moral responsibility to virtual humans [47]. Other untested boundary conditions include how transgression-related messages are framed [88], individual differences in speciesism against AI, a general bias favoring biological over artificial agents [89],

and broader environmental variables (e.g., virtual atmosphere, platform characteristics) and consumer traits (e.g., virtual interaction habits, platform usage patterns).

Third, the rapid pace of change in virtual influencer technology introduces a further limitation tied to timing rather than design. The low-agency perception of AI-driven virtual influencers observed in this study reflects the current generation of AI technology, which consumers generally perceive as reactive and rule-based, akin to a conversational model with a visual interface [65,66]; as large language models grow more autonomous, and as consumers' own understanding of what AI systems can do continues to evolve, this perceptual anchor may shift. As AI capabilities and public familiarity with them continue to evolve rapidly [90,91], the boundary condition identified in the present research, namely that AI-driven status attenuates blame, should be revisited periodically rather than treated as a stable feature of consumer psychology. Future studies could directly manipulate the perceived autonomy of the AI system (e.g., a simple scripted chatbot versus an agentic system with apparent initiative) to test whether the protective effect of AI-driven framing weakens as perceived AI autonomy increases, and should also attend to the broader ethical implications of increasingly capable AI-driven personas [92].

Author Contributions: Conceptualization, W.S. and S.W.; methodology, W.S., S.W. and Y.D.; software, Z.L.; validation, S.W., Z.L. and S.D.; formal analysis, W.S. and S.W.; investigation, W.S. and S.W.; resources, Y.D. and S.D.; data curation, Z.L.; writing—original draft preparation, W.S. and S.W.; writing—review and editing, S.D. and Y.D.; visualization, Z.L.; supervision, Y.D.; project administration, Y.D.; funding acquisition, Y.D. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by the National Natural Science Foundation of China (grant number 72202084) and the China Scholarship Council (CSC).

Institutional Review Board Statement: The study was conducted in accordance with the Declaration of Helsinki, and approved by the Ethics Committee of the Graduate School, Guizhou University of Finance and Economics (approval code GSIRB002, approved on 3 January 2026).

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

Data Availability Statement: The raw data supporting the conclusions of this article will be made available by the corresponding author on request.

Acknowledgments: The authors thank the editors and the anonymous reviewers for their constructive comments.

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

Appendix A. Cover Story for Manipulating Virtual Influencers' Appearance (Used in Study 1 and Study 2)

The character in the picture is a human-like (vs. anime-like) virtual influencer on a platform. The virtual influencers are not real people, but virtual images formed by artificial intelligence, face modeling, voice synthesis and other technologies. Their speech, action and image can be manipulated by complex algorithm technology. Virtual influencers are usually operated by a professional technical and commercial team. With the development of digital technology, more and more platforms have introduced virtual influencers. Virtual influencers also have their own personality settings. Like human influencers, they interact with followers and recommend products to them.

Appendix B. Transgression Scenarios

- (1) *Transgression based on the false endorsement (used in Study 1 and Study 2):* The influencer claimed to have selected an exclusive coffee cup with exquisite design for her fans and recommended them to buy it. However, after the fans bought the product, they found

that the coffee cup was not as exquisite as the influencer claimed, and the quality was poor, far from the fans' expectations.

- (2) *Transgression based on discrediting statements (used in Study 1)*: The influencer made discriminatory comments about disabled people during a live broadcast, which made followers feel very uncomfortable.
- (3) *Transgression based on discrediting statements (used in Study 3)*: Imagine you have followed Yuki for a long time. Recently, you noticed that Yuki regularly posted, liked and commented on Instagram posts related to antisemitism and supported for racism publicly. While some people pointed out that this is inappropriate behavior, Yuki replied that she is just joking and talking for fun, which have caused an outcry from her fans and the public.

Appendix C. Manipulation of Types of Virtual Influencers in Study 3

- (1) *Cover story for manipulating virtual influencers' appearance*. Yuki is a virtual influencer on Instagram with a "human-like" image (vs. cartoon image). Virtual influencers are computer-generated characters with social media presence. They are not real people, but usually run by a specialized technical and commercial team.
- (2) *Cover story for manipulating virtual influencers' driving mechanism*. Yuki is driven by a real person behind the scenes (vs. by AI), who wears a full set of body and facial motion capturing equipment, combined with voice synthesis technology (vs. formed by high-precision multimodal technology that restores movement transformations through the use of artificial intelligence and deep learning). Therefore, people visually see the virtual appearance as presented in the picture. Her actions are actually exhibited by the real person inside the holster (vs. actually ran through the complex algorithms) and the sounds she makes are also virtually processed. Recently, like human influencers, she is often active on various social media platforms, sharing her life and outfits, occasionally endorsing some products, and also liking and commenting to her fans.

References

- Hollebeek, L.D.; Macky, K. Digital content marketing's role in fostering consumer engagement, trust, and value: Framework, fundamental propositions, and implications. *J. Interact. Mark.* **2019**, *45*, 27–41. [\[CrossRef\]](#)
- Drenten, J.; Brooks, G. Celebrity 2.0: Lil Miquela and the rise of a virtual star system. *Fem. Media Stud.* **2020**, *20*, 1319–1323. [\[CrossRef\]](#)
- Molenaar, K. Discover The Top 12 Virtual Influencers for 2024—Listed and Ranked! Influencer Marketing Hub 29 March 2024. Available online: <https://influencermarketinghub.com/virtual-influencers/> (accessed on 15 December 2025).
- Deng, F.; Jiang, X. Effects of human versus virtual human influencers on the appearance anxiety of social media users. *J. Retail. Consum. Serv.* **2023**, *71*, 103233. [\[CrossRef\]](#)
- Lim, R.E.; Lee, S.Y. "You are a virtual influencer!": Understanding the impact of origin disclosure and emotional narratives on parasocial relationships and virtual influencer credibility. *Comput. Hum. Behav.* **2023**, *148*, 107897. [\[CrossRef\]](#)
- You, L.; Liu, F. From virtual voices to real impact: Authenticity, altruism, and egoism in social advocacy by human and virtual influencers. *Technol. Forecast. Soc. Change* **2024**, *207*, 123650. [\[CrossRef\]](#)
- Allal-Chérif, O.; Puertas, R.; Carracedo, P. Intelligent influencer marketing: How AI-powered virtual influencers outperform human influencers. *Technol. Forecast. Soc. Change* **2024**, *200*, 123113. [\[CrossRef\]](#)
- Ozdemir, O.; Kolfal, B.; Messinger, P.R.; Rizvi, S. Human or virtual: How influencer type shapes brand attitudes. *Comput. Hum. Behav.* **2023**, *145*, 107771. [\[CrossRef\]](#)
- Mouritzen, S.L.T.; Penttinen, V.; Pedersen, S. Virtual influencer marketing: The good, the bad and the unreal. *Eur. J. Mark.* **2023**. [\[CrossRef\]](#)
- Sands, S.; Ferraro, C.; Demsar, V.; Chandler, G. False idols: Unpacking the opportunities and challenges of falsity in the context of virtual influencers. *Bus. Horiz.* **2022**, *65*, 777–788. [\[CrossRef\]](#)
- Mrad, M.; Ramadan, Z.; Nasr, L.I. Computer-generated influencers: The rise of digital personalities. *Mark. Intell. Plan.* **2022**, *40*, 589–603. [\[CrossRef\]](#)

12. Mustak, M.; Salminen, J.; Mäntymäki, M.; Rahman, A.; Dwivedi, Y.K. Deepfakes: Deceptions, mitigations, and opportunities. *J. Bus. Res.* **2023**, *154*, 113368. [\[CrossRef\]](#)
13. Koebler, J. A Computer-Generated, Pro-Trump Instagram Model Said She Hacked Lil Miquela, Another CGI Instagram Model. *Vice*, 17 April 2018. Available online: https://www.vice.com/en_us/article/43bp79/lil-miquela-instagram-allegedly-hacked-bermuda (accessed on 15 December 2025).
14. Raphael, S. Meet Bermuda, The Most Controversial CGI Influencer on Instagram. *Refinery 29*, 20 December 2018. Available online: <https://www.refinery29.com/en-gb/bermuda-instagram-cgi-influencer> (accessed on 15 December 2025).
15. Thomas, V.L.; Fowler, K. Close encounters of the AI kind: Use of AI influencers as brand endorsers. *J. Advert.* **2021**, *50*, 11–25. [\[CrossRef\]](#)
16. Zhao, T.; Ran, Y.; Wu, B.; Wang, V.L.; Zhou, L.; Wang, C.L. Virtual versus human: Unraveling consumer reactions to service failures through influencer types. *J. Bus. Res.* **2024**, *178*, 114657. [\[CrossRef\]](#)
17. Park, S.; Sung, Y. The interplay between human likeness and agency on virtual influencer credibility. *Cyberpsychol. Behav. Soc. Netw.* **2023**, *26*, 764–771. [\[CrossRef\]](#) [\[PubMed\]](#)
18. Arsenyan, J.; Mirowska, A. Almost human? A comparative case study on the social media presence of virtual influencers. *Int. J. Hum.-Comput. Stud.* **2021**, *155*, 102694. [\[CrossRef\]](#)
19. Ma, Y.; Li, J. How humanlike is enough? Uncover the underlying mechanism of virtual influencer endorsement. *Comput. Hum. Behav. Artif. Hum.* **2024**, *2*, 100037. [\[CrossRef\]](#)
20. Yang, J.; Chuentawong, P.; Lee, H.; Chock, T.M. Anthropomorphism in CSR endorsement: A comparative study on humanlike vs. cartoonlike virtual influencers' climate change messaging. *J. Promot. Manag.* **2023**, *29*, 705–734. [\[CrossRef\]](#)
21. Jiang, K.; Zheng, J.; Luo, S. Green power of virtual influencer: The role of virtual influencer image, emotional appeal, and product involvement. *J. Retail. Consum. Serv.* **2024**, *77*, 103660. [\[CrossRef\]](#)
22. Kim, E.; Kim, D.; E, Z.; Shoenberger, H. The next hype in social media advertising: Examining virtual influencers' brand endorsement effectiveness. *Front. Psychol.* **2023**, *14*, 1089051. [\[CrossRef\]](#) [\[PubMed\]](#)
23. Appel, G.; Grewal, L.; Hadi, R.; Stephen, A.T. The future of social media inmarketing. *J. Acad. Mark. Sci.* **2020**, *48*, 79–95. [\[CrossRef\]](#) [\[PubMed\]](#)
24. Shin, M.; Song, S.W.; Chock, T.M. Uncanny valley effects on friendship decisions in virtual social networking service. *Cyberpsychol. Behav. Soc. Netw.* **2019**, *22*, 700–705. [\[CrossRef\]](#) [\[PubMed\]](#)
25. Yan, J.; Xia, S.; Jiang, A.; Lin, Z. The effect of different types of virtual influencers on consumers' emotional attachment. *J. Bus. Res.* **2024**, *177*, 114646. [\[CrossRef\]](#)
26. Qu, Y.; Baek, E. Let virtual creatures stay virtual: Tactics to increase trust in virtual influencers. *J. Res. Interact. Mark.* **2024**, *18*, 91–108. [\[CrossRef\]](#)
27. Crolic, C.; Thomaz, F.; Hadi, R.; Stephen, A.T. Blame the bot: Anthropomorphism and anger in customer–chatbot interactions. *J. Mark.* **2022**, *86*, 132–148. [\[CrossRef\]](#)
28. Johnson, R.D.; Marakas, G.M.; Palmer, J.W. Differential social attributions toward computing technology: An empirical investigation. *Int. J. Hum.-Comput. Stud.* **2006**, *64*, 446–460. [\[CrossRef\]](#)
29. Edwards, C.; Edwards, A.; Stoll, B.; Lin, X.; Massey, N. Evaluations of an artificial intelligence instructor's voice: Social Identity Theory in human-robot interactions. *Comput. Hum. Behav.* **2019**, *90*, 357–362. [\[CrossRef\]](#)
30. Quach, S.; Cheah, I.; Thaichon, P. The power of flattery: Enhancing prosocial behavior through virtual influencers. *Psychol. Mark.* **2024**, *41*, 1629–1648. [\[CrossRef\]](#)
31. Xie-Carson, L.; Magor, T.; Benckendorff, P.; Hughes, K. All hype or the real deal? Investigating user engagement with virtual influencers in tourism. *Tour. Manag.* **2023**, *99*, 104779. [\[CrossRef\]](#)
32. Xu, H.; Wu, Y. Do virtual endorsers have a country-of-origin effect? From the perspective of congruent explanations. *Technol. Forecast. Soc. Change* **2024**, *206*, 123530. [\[CrossRef\]](#)
33. Kuchmaner, C.A.; Wiggins, J.; Grimm, P.E. The role of network embeddedness and psychological ownership in consumer responses to brand transgressions. *J. Interact. Mark.* **2019**, *47*, 129–143. [\[CrossRef\]](#)
34. Lo, C.J.E.; Tsarenko, Y.; Tojib, D. Same scandal, different moral judgments: The effects of consumer-firm affiliation on weighting transgressor-related information and post-scandal patronage intentions. *Eur. J. Mark.* **2021**, *55*, 3162–3190. [\[CrossRef\]](#)
35. Theodorakis, I.G.; Painesis, G. Ad eroticism from a psychological distance perspective: Investigating its effects in light of consumers' sex, ethical judgments, and moral attentiveness. *J. Bus. Res.* **2022**, *142*, 524–539. [\[CrossRef\]](#)
36. Reinikainen, H.; Tan, T.M.; Luoma-Aho, V.; Salo, J. Making and breaking relationships on social media: The impacts of brand and influencer betrayals. *Technol. Forecast. Soc. Change* **2021**, *171*, 120990. [\[CrossRef\]](#)
37. Von Mettenheim, W.; Wiedmann, K.P. Influencer transgressions: The impacts on endorser and brand. *J. Media Econ.* **2023**, *35*, 28–62. [\[CrossRef\]](#)

38. Zheng, B. Streamers, Virtual Nationalists and Soft Power: China vs. the World? *China Studies*. 30 October 2020. Available online: <https://www.bakerinstitute.org/research/streamers-virtual-nationalists-and-soft-power-china-vs-world> (accessed on 15 December 2025).
39. Aw, E.C.X.; Labrecque, L.I. Celebrities as brand shields: The role of parasocial relationships in dampening negative consequences from brand transgressions. *J. Advert.* **2023**, *52*, 387–405. [\[CrossRef\]](#)
40. Wang, S.; Kim, K.J. Consumer response to negative celebrity publicity: The effects of moral reasoning strategies and fan identification. *J. Product. Brand Manag.* **2020**, *29*, 114–123. [\[CrossRef\]](#)
41. Yang, Y.; Hu, J. Self-diminishing effects of awe on consumer forgiveness in service encounters. *J. Retail. Consum. Serv.* **2021**, *60*, 102491. [\[CrossRef\]](#)
42. Kennedy, A.; Baxter, S.M.; Ilicic, J. Celebrity versus film persona endorsements: Examining the effect of celebrity transgressions on consumer judgments. *Psychol. Mark.* **2019**, *36*, 102–112. [\[CrossRef\]](#)
43. Rifon, N.J.; Jiang, M.; Wu, S. Consumer response to celebrity transgression: Investigating the effects of celebrity gender and past transgressive and philanthropic behaviors using real celebrities. *J. Product. Brand Manag.* **2023**, *32*, 517–529. [\[CrossRef\]](#)
44. Kim, Y.; Ho, T.H.; Tan, L.P.; Casidy, R. Factors influencing consumer forgiveness: A systematic literature review and directions for future research. *J. Serv. Theory Pract.* **2023**, *33*, 601–628. [\[CrossRef\]](#)
45. Sinha, J.; Lu, F.C. “I” value justice, but “we” value relationships: Self-construal effects on post-transgression consumer forgiveness. *J. Consum. Psychol.* **2016**, *26*, 265–274. [\[CrossRef\]](#)
46. Liu, F.; Lee, Y.H. Virtually responsible? Attribution of responsibility toward human vs. virtual influencers and the mediating role of mind perception. *J. Retail. Consum. Serv.* **2024**, *77*, 103685. [\[CrossRef\]](#)
47. Yan, X.; Mo, T.; Zhou, X. The influence of cultural differences between China and the West on moral responsibility judgments of virtual humans. *Acta Psychol. Sin.* **2024**, *56*, 161. [\[CrossRef\]](#)
48. Gray, H.M.; Gray, K.; Wegner, D.M. Dimensions of mind perception. *Science* **2007**, *315*, 619. [\[CrossRef\]](#) [\[PubMed\]](#)
49. Tzelios, K.; Williams, L.A.; Omerod, J.; Bliss-Moreau, E. Evidence of the unidimensional structure of mind perception. *Sci. Rep.* **2022**, *12*, 18978. [\[CrossRef\]](#) [\[PubMed\]](#)
50. Gray, K.; Wegner, D.M. Feeling robots and human zombies: Mind perception and the uncanny valley. *Cognition* **2012**, *125*, 125–130. [\[CrossRef\]](#) [\[PubMed\]](#)
51. Waytz, A.; Gray, K.; Epley, N.; Wegner, D.M. Causes and consequences of mind perception. *Trends Cogn. Sci.* **2010**, *14*, 383–388. [\[CrossRef\]](#) [\[PubMed\]](#)
52. Li, H.; Lei, Y.; Zhou, Q.; Yuan, H. Can you sense without being human? Comparing virtual and human influencers endorsement effectiveness. *J. Retail. Consum. Serv.* **2023**, *75*, 103456. [\[CrossRef\]](#)
53. Luo, X.M.; Tong, S.L.; Fang, Z.; Qu, Z. Frontiers: Machines vs. humans: The impact of artificial intelligence chatbot disclosure on customer purchases. *Mark. Sci.* **2019**, *38*, 937–947. [\[CrossRef\]](#)
54. Kevin, K.; Jaime, B. It feels, therefore it is: Associations between mind perception and mind ascription for social robots. *Comput. Hum. Behav.* **2024**, *153*, 108098. [\[CrossRef\]](#)
55. Küster, D.; Swiderska, A.; Gunkel, D. I saw it on YouTube! How online videos shape perceptions of mind, morality, and fears about robots. *New Media Soc.* **2021**, *23*, 3312–3331. [\[CrossRef\]](#)
56. Swiderska, A.; Küster, D. Avatars in pain: Visible harm enhances mind perception in humans and robots. *Perception* **2018**, *47*, 1139–1152. [\[CrossRef\]](#) [\[PubMed\]](#)
57. Park, S.; Kim, S.P.; Whang, M. Individual’s social perception of virtual avatars embodied with their habitual facial expressions and facial appearance. *Sensors* **2021**, *21*, 5986. [\[CrossRef\]](#) [\[PubMed\]](#)
58. Ham, J.; Li, S.; Looi, J.; Eastin, M.S. Virtual humans as social actors: Investigating user perceptions of virtual humans’ emotional expression on social media. *Comput. Hum. Behav.* **2024**, *155*, 108161. [\[CrossRef\]](#)
59. Yang, D.; Zhang, J.; Sun, Y.; Huang, Z. Showing usage behavior or not? The effect of virtual influencers’ product usage behavior on consumers. *J. Retail. Consum. Serv.* **2024**, *79*, 103859. [\[CrossRef\]](#)
60. Gray, K.; Young, L.; Waytz, A. Mind perception is the essence of morality. *Psychol. Inq.* **2012**, *23*, 101–124. [\[CrossRef\]](#) [\[PubMed\]](#)
61. Sullivan, Y.W.; Wamba, F.S. Moral judgments in the age of artificial intelligence. *J. Bus. Ethics* **2022**, *178*, 917–943. [\[CrossRef\]](#)
62. Arikan, E.; Altinigne, N.; Kuzgun, E.; Okan, M. May robots be held responsible for service failure and recovery? The role of robot service provider agents’ human-likeness. *J. Retail. Consum. Serv.* **2023**, *70*, 103175. [\[CrossRef\]](#)
63. Seymour, M.; Riemer, K.; Kay, J. Actors, avatars and agents: Potentials and implications of natural face technology for the creation of realistic visual presence. *J. Assoc. Inf. Syst.* **2018**, *19*, 4. [\[CrossRef\]](#)
64. Tang, M.T.; Zhu, V.L.; Popescu, V. Alterecho: Loose avatar-streamer coupling for expressive vtubing. In Proceedings of the 2021 IEEE International Symposium on Mixed and Augmented Reality (ISMAR), Bari, Italy, 4–8 October 2021; pp. 128–137. [\[CrossRef\]](#)

65. Hoyer, W.D.; Kroschke, M.; Schmitt, B.; Kraume, K.; Shankar, V. Transforming the customer experience through new technologies. *J. Interact. Mark.* **2020**, *51*, 57–71. [[CrossRef](#)]
66. Roman, D. Virtual Influencers: The Power of AI-Generated Influencer Marketing. Tech, Trends & Market Insights, 15 February 2024. Available online: <https://wearebrain.com/blog/ai-generated-virtual-influencers> (accessed on 15 December 2025).
67. Von der Pütten, A.M.; Krämer, N.C.; Gratch, J.; Kang, S.H. “It doesn’t matter what you are!” Explaining social effects of agents and avatars. *Comput. Hum. Behav.* **2010**, *26*, 1641–1650. [[CrossRef](#)]
68. Nowak, K.L.; Fox, J. Avatars and computer-mediated communication: A review of the definitions, uses, and effects of digital representations. *Rev. Commun. Res.* **2018**, *6*, 30–53. [[CrossRef](#)]
69. Seymour, M.; Yuan, L.; Riemer, K.; Dennis, A.R. Less Artificial, More Intelligent: Understanding Affinity, Trustworthiness, and Preference for Digital Humans. *Inf. Syst. Res.* **2024**, online. [[CrossRef](#)]
70. Pavone, G.; Meyer-Waarden, L.; Munzel, A. Rage against the machine: Experimental insights into customers’ negative emotional responses, attributions of responsibility, and coping strategies in artificial intelligence-based service failures. *J. Interact. Mark.* **2023**, *58*, 52–71. [[CrossRef](#)]
71. Garvey, A.M.; Kim, T.; Duhachek, A. Bad news? Send an AI. Good news? Send a human. *J. Mark.* **2023**, *87*, 10–25. [[CrossRef](#)]
72. Faul, F.; Erdfelder, E.; Buchner, A.; Lang, A.G. Statistical power analyses using G Power 3.1: Tests for correlation and regression analyses. *Behav. Res. Methods* **2009**, *41*, 1149–1160. [[CrossRef](#)] [[PubMed](#)]
73. Bigman, Y.E.; Gray, K. People are averse to machines making moral decisions. *Cognition* **2018**, *181*, 21–34. [[CrossRef](#)] [[PubMed](#)]
74. Hayes, A.F. *Introduction to Mediation, Moderation, and Conditional Process Analysis: A Regression-Based Approach*; Guilford Press: New York, NY, USA, 2013.
75. Yoo, J.W.; Park, J.; Park, H. How can I trust you if you’re fake? Understanding human-like virtual influencer credibility and the role of textual social cues. *J. Res. Interact. Mark.* **2025**, *19*, 730–748. [[CrossRef](#)]
76. Epley, N. A mind like mine: The exceptionally ordinary underpinnings of anthropomorphism. *J. Assoc. Consum. Res.* **2018**, *3*, 591–598. [[CrossRef](#)]
77. Zhou, X.; Yan, X.; Jiang, Y. Making sense? The sensory-specific nature of virtual influencer effectiveness. *J. Mark.* **2023**, *88*, 84–106. [[CrossRef](#)]
78. De Cicco, R.; Iacobucci, S.; Cannito, L.; Onesti, G.; Ceccato, I.; Palumbo, R. Virtual vs. human influencer: Effects on users’ perceptions and brand outcomes. *Technol. Soc.* **2024**, *77*, 102488. [[CrossRef](#)]
79. Yim, A.; Liska, L.; Mu, Y.; Thakkar, M. I feel no empathy toward AI: The effects of AI vs human spokespersons on advertisement effectiveness. *J. Res. Interact. Mark.* **2026**, *20*, 108–122. [[CrossRef](#)]
80. Sun, C.; Ye, C.; Li, C.; Liu, Y. Virtual ideality vs. virtual authenticity: Exploring the role of social signals in interactive marketing. *J. Res. Interact. Mark.* **2024**, *18*, 430–445. [[CrossRef](#)]
81. Chopdar, P.K.; Paul, J. The impact of brand transparency of food delivery apps in interactive brand communication. *J. Res. Interact. Mark.* **2024**, *18*, 238–256. [[CrossRef](#)]
82. Shen, P.; Nie, X.; Tong, C. Does disclosing commercial intention benefit brands? Mediating role of perceived manipulative intent and perceived authenticity in influencer hidden advertising. *J. Res. Interact. Mark.* **2025**, *19*, 673–693. [[CrossRef](#)]
83. Xin, B.; Hao, Y.; Xie, L. Virtual influencers and corporate reputation: From marketing game to empirical analysis. *J. Res. Interact. Mark.* **2024**, *18*, 759–786. [[CrossRef](#)]
84. Fărte, G.; Obadă, D.R.; Gherguț-Babii, A.; Dabija, D. Building corporate immunity: How do companies increase their resilience to negative information in the environment of fake news? *J. Res. Interact. Mark.* **2025**, *19*, 1260–1277. [[CrossRef](#)]
85. Yang, J.; Basile, K.; Zhao, X. Examining CSR communication on social media during a victim crisis: A machine learning based text analytics approach. *J. Res. Interact. Mark.* **2025**, *19*, 840–860. [[CrossRef](#)]
86. Shao, Z. Understanding the switching intention to virtual streamers in live streaming commerce: Innovation resistances, shopping motivations and personalities. *J. Res. Interact. Mark.* **2025**, *19*, 333–357. [[CrossRef](#)]
87. Wei, Y.; Syahrivar, J.; Simay, A.E. Unveiling the influence of anthropomorphic chatbots on consumer behavioral intentions: Evidence from China and Indonesia. *J. Res. Interact. Mark.* **2025**, *19*, 132–157. [[CrossRef](#)]
88. Zheng, X.; Cui, C.; Zhang, C.; Li, D. Who says what? How message appeals shape virtual- versus human-influencers’ impact on consumer engagement. *J. Res. Interact. Mark.* **2026**, *20*, 199–214. [[CrossRef](#)]
89. Cheng, J.; Wang, J. Influencer-product attractiveness transference in interactive fashion marketing: The moderated moderating effect of speciesism against AI. *J. Res. Interact. Mark.* **2025**, *19*, 712–729. [[CrossRef](#)]

90. Peltier, J.W.; Dahl, A.J.; Schibrowsky, J.A. Artificial intelligence in interactive marketing: A conceptual framework and research agenda. *J. Res. Interact. Mark.* **2024**, *18*, 54–90. [[CrossRef](#)]
91. Wang, C.L. Editorial: The changing landscape of marketing research in the AI era: Prospects and challenges. *J. Res. Interact. Mark.* **2026**, *20*, 1–10. [[CrossRef](#)]
92. Labrecque, L.I.; Peña, P.Y.; Leonard, H.; Leger, R. Not all sunshine and rainbows: Exploring the dark side of AI in interactive marketing. *J. Res. Interact. Mark.* **2024**, *18*, 970–999. [[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.