

2026



Tool or Replacement?

Why Consumers Accept Some AI
Marketing Functions and Reject Others

1. Executive Summary

Marketing and sales lead all business functions in AI adoption (McKinsey, 2025), and that adoption is accelerating. Yet consumer sentiment is moving in the opposite direction. The IAB and Sonata Insights (2024, 2026) found a 37-point gap between executive confidence in AI-driven advertising and consumer comfort with it, and that gap widened over the two-year period. What has been missing from this conversation is a controlled comparison across functions. Practitioners know consumers resist some AI applications, but without side-by-side data, they cannot tell whether the resistance is about AI itself or about specific uses of it.

This paper provides that comparison. One hundred and fifty U.S. adults evaluated four consumer-facing AI marketing functions under identical disclosure conditions in a within-subjects survey: AI creative advertising, AI personalized content delivery, AI customer service chatbots, and AI virtual brand spokespeople. Each participant rated all four on brand attitude, perceived authenticity, and purchase intention. Practitioner evidence, drawn from a formal interview and field observation in the gaming sector alongside solicited industry responses, supplemented the consumer data.

The results show a two-tier structure. Personalized content delivery and customer service chatbots scored at or above the scale midpoint on all three measures. Creative advertising and virtual spokespeople scored below it. All cross-tier gaps were significant at $p < .001$. Within each tier, differences were not. The same disclosure language produced sharply different consumer responses depending on which function the AI occupied. The dividing line is source type: consumers evaluate AI differently when they

perceive it as a functional system than when they perceive it as occupying a role previously held by a human. Within each tier, the two functions arrive at similar scores through different mechanisms. Section 9 develops these distinctions in full.

Three cross-cutting principles emerge for practitioners making deployment decisions. First, replacement framing is the primary risk variable: if a reasonable consumer could read the deployment as replacing an identifiable human role, the function enters a higher risk tier. Second, audience composition shapes bottom-tier outcomes: aggregate sentiment scores mask how strongly audience composition shapes bottom-tier outcomes for creative advertising and virtual spokespeople. Third, disclosure should be treated as a design constraint: with the EU AI Act's consumer-facing obligations taking effect in August 2026, a function's position in the two-tier structure should drive both how much to invest in it and how to execute the disclosure.

The study is based on a single consumer sample (N = 150) recruited through Prolific, with practitioner evidence concentrated in the gaming sector. The findings establish a controlled baseline for a cross-function comparison that did not previously exist. Chapter 10 translates this baseline into function-specific deployment checklists.

2. Introduction: The Gap Nobody's Comparing

Practitioners are deploying AI across marketing and communications at a pace that has no precedent. McKinsey's 2025 State of AI research, drawing on data from thousands of business leaders globally, identified marketing and sales as one of the two business functions most consistently reporting AI use across eight consecutive years of research. Among PR practitioners specifically, the numbers are similarly striking. The USC Annenberg Center for Public Relations' 2025 Global Communication Report found that 60% of PR professionals surveyed believe AI will have a positive effect on the PR industry. Social media content (43%), press material development (36%), marketing communications (30%), and creative services (26%) are all functions where respondents reported current organizational use. Muck Rack's State of PR 2025 survey found that 77% of PR practitioners are already using generative AI in their daily work.

These figures describe an industry that has moved past the question of whether to use AI and is now making operational choices about where and how. The decisions being made today are not experimental. They shape what consumers see and how brands address consumers directly.

Consumer sentiment is moving in the opposite direction.

SurveyMonkey's AI Sentiment Study, conducted among 2,901 U.S. adults, found that the share of Americans expressing concern about AI rose from 27% to 33% in a single year. A Gartner survey published in 2026, drawing on responses from 335 U.S. consumers, found that 78% identified explicit labeling of AI-generated content as either "very important" or "the most important factor" in maintaining trust. A 2024 Salesforce

survey found that only 49% of consumers believe companies use their data in ways that benefit them, down from 60% in 2022. Across multiple independent sources, the consumer trust picture is deteriorating even as deployment accelerates.

The most direct evidence comes from the Interactive Advertising Bureau and Sonata Insights. Their 2026 research found that 82% of advertising executives believe Gen Z and millennial consumers feel positively about AI-generated advertising. Only 45% of those consumers actually do. That 37-point gap represents a systematic miscalibration at the center of one of the industry's most consequential technology decisions. The same research team documented the same gap at 32 points just two years earlier. It is widening, not narrowing.

You might expect that as AI use in marketing has grown, practitioners would be refining their sense of how consumers respond. The data says the opposite is true.

What the existing research does not tell you

The practitioner-consumer gap documented above is real, but the research that captures it has a structural limitation. Most AI marketing research examines one function at a time. Studies of consumer response to AI-generated advertising are not compared against studies of consumer response to AI customer service chatbots. Research into virtual influencers sits in a separate literature from research into personalized content delivery. Practitioners making decisions across multiple AI marketing functions are working from a collection of isolated findings that cannot tell them whether consumer response varies by function.

Existing studies examine these functions in isolation. No single study has placed all four in front of the same consumer under matched conditions. That is the gap this white paper is designed to fill.

What this white paper does

This research examines consumer sentiment toward four AI marketing functions: AI creative advertising, AI personalized content delivery, AI customer service chatbots, and AI virtual brand spokespeople. These four functions represent prominent consumer-facing AI marketing applications that are currently in active use or active evaluation across the industry.

The core method is a within-subjects consumer survey. One hundred and fifty U.S. adults each evaluated all four functions under identical conditions, using consistent scenario formats and consistent measurement scales. That design makes direct cross-function comparison possible in a way that single-function studies and heterogeneous literature reviews cannot replicate. Each participant rated every function on brand attitude, perceived authenticity, and purchase intention, giving the comparison three independent dimensions.

A second research stream adds practitioner perspective. Evidence drawn from a formal interview, field observation, and solicited industry responses documents how deployment decisions are actually made across these four functions and what factors practitioners weigh. The consumer and practitioner findings are interpreted together in Section 9 and translated into deployment recommendations in Section 10.

This study represents, to the author's knowledge, the first controlled within-subjects cross-function comparison of consumer sentiment in this area.

3. Theoretical Foundations

Consumer sentiment toward AI in marketing does not follow a single pattern. It shifts depending on the function, the context, and what the consumer thinks the AI is trying to do. Five theoretical frameworks help explain why. Each maps directly onto one or more of the four functions examined in this research.

3.1 How Consumers Process AI Persuasion

Consumers are not passive recipients of marketing messages. Over time, they develop personal knowledge about the goals and tactics of persuasion agents, and they draw on that knowledge whenever they encounter a message they suspect is commercially motivated. This is the foundation of the Persuasion Knowledge Model (PKM), introduced by Friestad and Wright (1994). When a consumer recognizes a situation as a persuasion attempt, this can activate persuasion knowledge and heighten scrutiny. The consumer may attend not only to the content of the message but also to the motives and tactics behind it. Importantly, the model treats coping as a neutral process: the goal is not necessarily to resist, but to maintain control over the outcome of the interaction.

Watson, Valsesia, and Segal (2024) extend this framework to AI marketing contexts, conceptualizing AI as a persuasion agent and reviewing the conditions under which consumers respond to it. They identify two preconditions for the PKM lens to become relevant: the consumer must recognize that the agent they are interacting with is AI, and the consumer must perceive that the agent carries some persuasive intent. When both conditions are present, AI-related persuasion knowledge becomes active and the consumer begins evaluating the AI's role and motives. Whether this produces

resistance or acceptance depends on additional factors, including the properties of the AI itself, individual differences in AI attitudes, and the context of the interaction.

Qiu, Wang, Zeng, and Cong (2025) provide empirical evidence of this mechanism in a marketing context. Across two experiments on AI-disclosed cause-related marketing, they find that AI disclosure increases consumer skepticism, which reduces advertisement attitude, which in turn reduces purchase intention. The inhibitory effect of AI disclosure was significant in cause-related marketing contexts but not in standard commercial contexts, and a second study confirmed the sequential mediation pathway within cause-related marketing specifically. The effect was stronger among consumers who already held negative views of AI. Their findings extend the PKM into generative AI contexts and show that the inhibitory effect of AI disclosure is not uniform, varying across marketing contexts and with consumers' preexisting attitudes toward AI.

The PKM raises a specific question for this research: do the four functions differ in how readily they activate persuasion knowledge, and does that variation predict differences in consumer sentiment? The answer depends on whether both preconditions Watson et al. identify are present in each function: whether the consumer recognizes the agent as AI, and whether they perceive it as carrying persuasive intent.

3.2 Algorithm Aversion and Appreciation

Consumers do not uniformly distrust AI, nor do they uniformly accept it. Research shows they will defer to algorithmic judgment in some contexts and strongly resist it in others. Two foundational studies establish this pattern. Dietvorst, Simmons, and Massey (2015) demonstrated that people lose confidence in algorithmic forecasters

more quickly than in human forecasters after seeing them make the same mistake. Even when participants watched an algorithm outperform a human, they became less willing to use it after seeing it err. People tolerate human error in ways they do not extend to machines. Logg, Minson, and Moore (2019) showed the reverse is also possible: in a range of estimation tasks, people actually preferred algorithmic advice over human advice when no prior performance information made the algorithm's fallibility salient. The two findings together establish that algorithm aversion and algorithm appreciation are both real, and the question is what determines which response dominates.

Qin et al. (2025) synthesize this literature in a meta-analysis of 442 effect sizes from 163 studies ($N = 82,078$), and propose the Capability-Personalization Framework to explain the divergence. The framework identifies two dimensions consumers evaluate when deciding whether to rely on AI versus a human in a given context: the perceived capability of the AI relative to a human, and the perceived necessity for personalization in that context. AI appreciation occurs only when both conditions are met simultaneously: AI is perceived as more capable than a human, and personalization is perceived as unnecessary in the decision context. In the other three quadrants of the framework, where one or both conditions are absent, AI aversion occurs. The meta-analytic evidence supports this: AI appreciation produced a mean effect size of $d = 0.27$, while AI aversion produced $d = -0.50$, indicating that the negative preference for humans over AI was larger in magnitude than the positive preference for AI over humans in their meta-analytic estimates.

The Capability-Personalization Framework raises a pointed question for this research: do the four functions differ systematically in how consumers assess AI capability and the necessity for personalization? Personalized content is a particularly interesting test case. By definition it is a context where personalization is central, yet consumer demand for recommendation-based features is well documented. Whether consumers perceive data-driven product matching as demanding the kind of human judgment the framework treats as necessary, or as a computational task AI handles adequately, is an empirical question this study is positioned to answer.

3.3 Source Effects in Non-Human Communication

When a consumer evaluates a message, they evaluate both what is being said and who or what is saying it. Source attribution shapes how content is processed. Source credibility research has long established that who delivers a message affects how it is received, independent of the message content itself (Hovland, Janis, & Kelley, 1953). What happens when the source is not a person?

Sundar and Nass (2001) address this question directly by proposing a typology of sources applicable to online communication. They distinguish three categories. Visible sources are human gatekeepers or presenters perceived as delivering the message. Technological sources are media or channels perceived by receivers as autonomous originators of content. Receiver sources encompass the audience or the self as content selectors. In an experiment where participants read identical news stories attributed to one of four sources, news editors, a computer terminal, other users, or themselves, attribution to different source types produced significant variation in how participants rated the stories on liking, quality, and representativeness, though not in perceived

credibility, a pattern that underscores that source effects do not uniformly affect all dimensions of content evaluation.

The conceptual contribution of the typology is its most direct relevance to AI marketing. It establishes that the category of source, not just the content itself, is an independent variable in how receivers evaluate communication. A visible human source carries associations of editorial judgment and human agency. A technological source is processed as autonomous but not independent, meaning receivers understand it is programmed while still responding to it as if it operates on its own terms. This distinction between source types anticipates a question central to AI marketing: when a brand deploys AI as a visible, anthropomorphized entity rather than as a background process, how does the consumer's source attribution change, and what does that change do to their evaluation of the message?

3.4 AI Disclosure Effects

Recent experimental work shows that AI disclosure changes trust, credibility, and evaluative responses to marketing content, though the direction and magnitude of the effect depend on context. Four studies document how this works.

Bui (2025) tested this directly in an advertising context. Using 358 Prolific participants exposed to AI-generated advertisements, the study found that when the AI origin of an advertisement was disclosed, perceived advertising value and purchase intentions were lower than when the AI source was concealed. The relationship was mediated by advertising credibility: disclosure reduced how credible consumers found the ad, and lower credibility in turn reduced both value perceptions and purchase intent. Consumer

attitudes toward AI moderated the effect, such that more favorable AI attitudes weakened the negative impact of disclosure. The finding establishes that in commercial advertising contexts, transparency about AI involvement carries a cost to credibility that flows through to behavioral outcomes.

Koning and Voorveld (2025) extend this analysis to organizational trust. In a between-subjects experiment ($N = 304$), participants exposed to a GAI disclosure showed higher conceptual AI knowledge and higher attitudinal persuasion knowledge, and these activations were associated with lower trust in both the advertisement and the organization. The study also reported a smaller positive pathway, showing that disclosure can have trust-enhancing effects under some evaluative conditions. The overall pattern was therefore negative but not uniform. The authors conclude that the direction of the effect depends on which type of persuasion knowledge disclosure activates. When disclosure raises conceptual AI knowledge, awareness that AI was involved, consumers may begin to question whether using AI in advertising is appropriate. This increase in attitudinal persuasion knowledge is associated with lower trust in both the ad and the organization. When disclosure is processed as a sign of brand transparency rather than a manipulation cue, it can modestly increase perceived appropriateness and trust.

Wortel, Vanwesenbeeck, and Tomas (2024) examined AI disclosures in Instagram advertising and found a pattern that did not map cleanly onto the negative-disclosure hypothesis. Their study found that disclosure may have been processed partly as a transparency cue, even though advertising attitudes still declined. Contrary to the study's initial predictions, disclosures did not increase inferences of manipulative intent.

In fact, those inferences were lower in the disclosure condition. Even among high-AI-aversion consumers, the disclosure effect emerged in brand attitude rather than source credibility, which remained statistically unchanged. The study's findings suggest that negative AI ad evaluations are not primarily explained by a manipulation inference pathway, though the alternative mechanisms, which may include perceived authenticity, humanness, or creative effort, were not directly tested in this study.

Grigsby, Michelsen, and Zamudio (2025) add an important boundary condition. In a series of three experiments focused on service advertising, they found that the negative effects of AI disclosure on trust and ad attitudes were not uniform across ad designs. When an AI-disclosed advertisement focused on intangible service attributes, such as the image of a service provider, it generated significantly less trust than an equivalent ad focusing on tangible attributes like equipment or environment. Critically, trust and ad attitudes could be largely restored when AI was used to generate only the tangible elements of the ad while the intangible elements remained human-produced. The negative effect of disclosure is not fixed: it concentrates in the ad elements consumers use to infer human judgment, warmth, and relational intent.

Taken together, these four studies describe a disclosure environment that is neither uniformly negative nor straightforwardly managed by transparency alone. Disclosure activates scrutiny, and that scrutiny concentrates in the elements consumers rely on to judge trustworthiness. The harm from disclosure concentrates where consumers expect human judgment and relational authenticity.

3.5 Psychological Distance and AI Function Type

The four theoretical frameworks covered so far each illuminate a different dimension of how consumers respond to AI in marketing contexts. Persuasion knowledge explains why consumers scrutinize AI more closely when they recognize it as a persuasion agent. The capability-personalization framework explains why consumer response varies by function: appreciation emerges when AI is seen as more capable than humans and personalization is perceived as unnecessary, and aversion occurs otherwise. Source attribution theory explains why the category of source, not just content, shapes how messages are evaluated. Disclosure effects show that making AI involvement visible changes consumer trust, with the direction depending on what kind of knowledge the disclosure activates. A fifth framework adds a dimension that the others do not address directly: the degree to which a consumer feels psychologically close to, or distant from, the AI interaction itself.

Construal Level Theory (CLT), developed by Trope and Liberman (2010), proposes that psychological distance shapes how people mentally represent objects and events. The theory defines psychological distance as a subjective experience of how close or far something is from the self in the here and now, with four dimensions: temporal distance, spatial distance, social distance, and hypotheticality. A core principle of CLT is that as psychological distance increases, mental representations become more abstract, high-level, and focused on central features. Conversely, as psychological distance decreases and an object or event feels proximate, representations become more concrete, low-level, and focused on contextual details. Prediction, preference, and evaluation are all shaped by this construal process.

CLT has not been applied to AI marketing contexts, and the extension here is a theoretical inference. The framework raises a question relevant to cross-function comparison: if the four AI marketing functions differ in how psychologically close or distant they feel to the consumer, do they also produce different modes of evaluation, shifting from concrete, feasibility-based judgment at the proximate end toward abstract, value-laden judgment at the distal end?

4. Four Functions in Context

The theoretical frameworks in Section 3 describe general mechanisms that shape consumer responses to AI. This section applies those frameworks to each of the four marketing functions examined in this research. For each function, the goal is to characterize the existing state of knowledge: what practitioners are doing, what consumers are reporting, and where the two diverge. This secondary literature is the foundation on which the primary research in Sections 7 and 8 builds.

4.1 AI Creative Advertising

AI creative advertising is the most publicly visible of the four functions and the one for which the gap between practitioner behavior and consumer sentiment is best documented. The Interactive Advertising Bureau and Sonata Insights report in their 2026 research that 83% of advertising executives now report deploying AI in the creative process, up from 60% in 2024. Over the same period, consumer sentiment moved in the opposite direction. The share of Gen Z and millennial consumers expressing negative sentiment toward AI-generated advertising grew by 12 percentage points. The percentage of consumers describing AI-using brands as innovative dropped from 30% to 23%, while the percentage of advertising executives holding the same belief rose from 40% to 49%. The two groups are diverging in attitude and moving in opposite directions at the same time.

The cognitive mechanism behind consumer resistance has been documented experimentally. Bui (2025) found that when the AI origin of an advertisement was disclosed to participants, they evaluated the ad's value and their purchase intentions

more negatively than when the same ad was presented without disclosure. The mediating variable was advertising credibility: disclosure reduced perceived credibility, which in turn reduced both value perceptions and purchase intent. The study recruited 358 participants via Prolific and used actual AI-generated stimuli, making the findings directly relevant to the deployment environment practitioners now operate in. Consumer attitudes toward AI moderated the effect, suggesting that the credibility penalty from disclosure concentrates among consumers who are already skeptical of AI, rather than falling evenly across the sample.

Koning and Voorveld (2025) extend the picture from advertisement-level outcomes to organizational trust. In their experiment ($N = 304$), AI disclosure in a generative AI advertisement was associated with lower trust in both the advertisement and the organization. The study also identified a smaller positive pathway, where disclosure increased perceived appropriateness and modestly raised trust, consistent with the mixed findings in the broader disclosure literature. The net direction was negative, and the implication extends beyond individual campaign performance: a consumer who trusts the brand less after seeing an AI-disclosed ad may carry that evaluation into subsequent interactions.

Chen, Wang, Rao Hill, and Li (2024) show that the picture is more nuanced than a simple anti-AI conclusion. Across four experiments, they find that consumer attitudes toward AI-generated advertising depend substantially on the type of appeal the ad uses. Consumers respond more positively to AI-generated ads that use agentic appeals, which emphasize individual task performance and self-improvement, than to AI-generated ads with communal appeals, which emphasize social relationships and

warmth. The mechanism is self-efficacy: agentic AI ads raise task self-efficacy, while communal content is evaluated more favorably when created by humans, because human-created ads better support social self-efficacy. Assigning a social role to the AI generator, framing it as a servant rather than an autonomous creator, can partially mitigate the negative effect of communal appeal AI ads. The study was conducted with Chinese consumers, and the authors note that cultural factors may limit generalizability to Western markets. The finding nonetheless suggests that the negative consumer response to AI creative advertising is not monolithic: it varies with what the ad is trying to do emotionally and relationally.

NielsenIQ's 2024 research into consumer attitudes toward AI-generated advertising identified three patterns worth noting. First, consumers often identify AI-generated advertisements on their own, describing them as annoying, boring, or confusing, and rating them as less engaging than traditional ads. Second, the study found evidence of a negative halo effect: AI-generated advertising reduced evaluation of the ad itself and dragged down broader brand perception as well. Third, neurological measurement using EEG found that AI-generated ads produced weaker memory activation than human-created equivalents, suggesting they may be harder to recall and less likely to motivate behavior. These findings are from an industry research report rather than a peer-reviewed study, but they point to consumer resistance that operates independently of explicit disclosure.

Taken together, the existing literature on AI creative advertising describes a function under active deployment pressure and facing documented consumer resistance. The resistance is not uniform: it varies with disclosure, with the type of appeal, and with prior

AI attitudes. But the directional finding is consistent across multiple independent studies and industry datasets. Practitioners are scaling AI use in creative at the same time consumers are becoming less receptive to it, and the gap between these two trends has been widening for two consecutive years.

4.2 AI Personalized Content Delivery

AI personalized content delivery occupies a distinctive position in the consumer landscape. It is the function consumers most directly benefit from, the one that generates the most data-related anxiety, and the one where the line between helpful and invasive is most actively contested. Understanding consumer response to this function requires holding two simultaneous truths: personalization is valued, and the data practices that enable it are not.

Salesforce's 2024 State of the Connected Customer report, drawing on a survey of more than 15,000 consumers worldwide, documents this tension directly. Seventy-three percent of consumers report that brands now treat them as unique individuals, suggesting that personalization at scale has become a baseline expectation rather than a novelty. Yet only 49% believe companies use their data in ways that benefit them, down from 60% in 2022, and 71% report being increasingly protective of their personal information. Consumers are receiving personalization and becoming more suspicious of it at the same time.

Hardcastle, Vorster, and Brown (2025) investigate the mechanisms behind this ambivalence in a qualitative study combining semi-structured phenomenological interviews with Customer Journey Mapping. Their findings show that consumers

navigate a persistent tension between two competing desires: the convenience of tailored recommendations and the discomfort of the surveillance those recommendations require. Participants in the study wanted personalized experiences but resented the extent of data collection that powered them. When personalization worked well, it was taken for granted. When it failed, through irrelevant or repetitive recommendations, the failure highlighted the data exchange that consumers had implicitly accepted, making it feel like a poor trade. The study also found that consumers are increasingly aware of dataveillance but remain uncertain about how recommendation algorithms actually function, which compounds mistrust: they know their data is being used but cannot evaluate whether it is being used fairly.

Hardcastle et al. situate this tension within a broader ethical concern about consumer autonomy. AI-driven recommendation systems narrow the options consumers see, which can streamline decision-making but can also limit exposure to alternatives in ways consumers do not consciously register. The study describes participants feeling “pigeonholed” by systems that kept surfacing the same categories of products, creating a feedback loop that reinforced existing preferences rather than expanding them. This dimension of personalization, its tendency to constrain rather than expand consumer choice, is not captured in satisfaction ratings and does not appear in quantitative attitude measures, but it shapes the qualitative character of how consumers relate to these systems over time.

The Qin et al. (2025) capability-personalization framework introduced in Section 3.2 would predict ambivalence for this function specifically. Personalized content is, by definition, a context where the perceived necessity for personalization is high, which the

framework suggests should increase the likelihood of AI aversion. Yet consumer demand for personalized experiences is documented and real, and practitioner investment in this function continues to grow. The resolution to this apparent contradiction may lie in how personalization is understood and evaluated by consumers. Hardcastle et al.'s findings suggest that personalization experienced as service enhancement produces different outcomes than personalization perceived as commercial extraction or surveillance. The framing and transparency of AI involvement in personalization may matter as much as the personalization itself.

Compared to AI creative advertising, personalized content operates closer to the consumer's own data and stated preferences. The AI is not imposing a creative vision. It reflects back what the consumer has already demonstrated. This distinction may help explain why the function generates less outright resistance than advertising, even as the underlying data practices raise parallel concerns. Whether consumers perceive the AI as working for them or on them appears to be a key variable, and it is one that framing, disclosure, and system behavior all shape.

4.3 AI Customer Service Chatbots

Customer service chatbots are the most established of the four functions. Enterprise deployment is mature, consumer exposure is widespread, and the research base is comparatively deeper than for the other functions examined here. That maturity does not mean consumer sentiment is uniformly positive. It means the shape of consumer resistance is better understood.

The most consequential study on chatbot disclosure and consumer behavior is Luo, Tong, Fang, and Qu (2019), a field experiment involving more than 6,200 customers randomized to receive outbound sales calls from either AI chatbots or human agents under different disclosure conditions. The core finding is striking: undisclosed chatbots performed as effectively as proficient human workers in generating customer purchases, and four times more effectively than inexperienced workers. But when chatbot identity was disclosed before the conversation began, purchase rates fell by more than 79.7%. Customers who knew they were speaking to a bot were shorter in their responses, more likely to hang up, and less likely to buy. The mechanism was not deception or manipulation concerns. Customers rated the disclosed bot as less knowledgeable and less empathetic than a human agent, even when its objective performance was equivalent. Two factors mitigated the effect: disclosing after rather than before the conversation, and prior AI experience among customers.

The Luo et al. study was conducted in a structured outbound sales context with Chinese consumers, and the authors explicitly note that their setting, limited to a restricted two-way information exchange, may not generalize to more dynamic inbound customer service interactions. The 79.7% figure should therefore be understood as specific to that context rather than as a universal estimate of disclosure cost. The directional finding, that disclosure activates resistance, is consistent with the broader disclosure literature reviewed in Section 3.4.

Industry data suggest that the practitioner picture is more complex than simple deployment enthusiasm. Gartner's 2024 survey of 5,728 U.S. consumers found that 64% would prefer companies avoid AI for customer service altogether, and that 53%

would consider switching to a competitor if they learned a company was using AI for service, with difficulty reaching a human agent cited as the primary concern. These figures indicate that consumer resistance to chatbots is not marginal and is concentrated in interactions where human judgment or empathy is perceived as necessary. A separate Gartner projection from 2022 had forecast that chatbots would become the primary customer service channel for roughly a quarter of organizations by 2027, a forecast that reflects the scale of enterprise investment in this direction regardless of consumer preference.

Reich, Kaju, and Maglio (2023) offer a useful complementary frame for understanding why chatbot resistance may persist even under conditions of high exposure. Their research finds that consumers avoid algorithmic advice partly on the faulty assumption that algorithms, unlike humans, cannot learn from mistakes. When consumers are shown that an algorithm has the capacity to improve, aversion decreases. For customer service chatbots, where failures are immediately visible, the implication is that making the system's learning capacity perceptible may be one lever for reducing resistance over time. Familiarity alone, as the Luo et al. findings suggest, does not automatically translate into acceptance.

The overall picture for AI customer service chatbots is one of functional acceptance and conditional tolerance. Consumers often end up using chatbots in part because they are available and positioned as faster than human alternatives, not necessarily because they prefer them. Resistance activates most reliably when the interaction involves emotional complexity, service recovery, or high-stakes contexts, where the Qin et al. capability-personalization framework would predict aversion because the perceived

necessity for human judgment is high. For routine, low-stakes queries, the same framework would predict the opposite: AI is seen as capable, personalization is unnecessary, and acceptance follows.

4.4 AI Virtual Brand Spokespeople

AI virtual brand spokespeople are the most experimental of the four functions and the clearest case where market growth and consumer sentiment move in opposite directions. The virtual influencer market was valued at \$6.33 billion in 2024 and is projected to grow at a compound annual growth rate of 38.4% through 2033, according to Straits Research market analysis. Yet consumer data consistently shows negative majorities. Sprout Social's Q3 2025 Pulse Survey found that 46% of consumers are uncomfortable with brands using AI influencers, with only 23% expressing comfort. A 2022 Statista survey of U.S. consumers found that 65% reported being unlikely to purchase products promoted by virtual influencers. These are industry and market research reports rather than peer-reviewed studies, and their methodologies vary, but the directional consistency across independent sources is notable.

The academic literature on virtual influencers reveals a more nuanced picture than the aggregate negative sentiment suggests. Franke, Groeppel-Klein, and Müller (2023) find that consumer responses to virtual influencers depend substantially on whether the virtual nature of the influencer is disclosed. When a virtual influencer is not labeled as such, consumers who later realize its non-human nature experience a stronger uncanny valley effect, characterized by feelings of eeriness and discomfort, than those who were informed from the outset. Franke et al. suggest that early disclosure can reduce the negative reactions associated with later realizing that the influencer is artificial, a finding

that complicates the assumption that concealing AI identity is always the safer commercial strategy. The study also shows that the effectiveness of virtual influencers varies with product category, suggesting that context shapes consumer tolerance for AI-generated personas.

Kim and Wang (2024) compare human-like virtual influencers, anime-like virtual influencers, and human influencers across for-profit and not-for-profit messaging contexts. Their findings are context-dependent in an important way: human-like virtual influencers can be as effective as human influencers in not-for-profit messaging, where perceived authenticity and source credibility are comparatively higher. When the messaging shifts to for-profit motives, that advantage disappears, and human-like virtual influencers perform closer to anime-like virtual influencers, which consistently show lower effectiveness. Their results show that perceived authenticity mediated the relationship between influencer type and marketing effectiveness in the not-for-profit condition. The mediating role of source credibility was not statistically supported in their tests, despite the study's emphasis on credibility as a theoretical mechanism.

Koles, Audrezet, Moulard, Ameen, and McKenna (2024) take a qualitative approach, drawing on interviews with 64 consumers and 11 industry experts to examine how authenticity manifests in virtual influencer contexts. Using an Entity-Referent Correspondence Framework of Authenticity, they identify three types of authenticity that apply to virtual influencers: true-to-ideal, which concerns whether the influencer lives up to a consistent set of values, true-to-fact, which concerns whether the representation is transparent and non-deceptive, and true-to-self, which concerns consistency of identity and expression over time. Their findings suggest that consumers do not simply reject

virtual influencers as inauthentic but apply a more differentiated evaluation framework that depends on how the influencer is positioned and what kind of authenticity the consumer prioritizes.

Lee, Shin, Yang, and Chock (2025) identify individual differences as a key moderator. Their study finds that machine heuristic, defined as the degree to which a consumer believes machines are objective, accurate, and unbiased, moderates how consumers evaluate virtual influencers. Consumers with a stronger machine heuristic rate virtual influencers as more authentic, suggesting that some consumers evaluate AI-generated personas using more machine-positive criteria than they apply to human influencers. This suggests that the negative aggregate sentiment toward virtual spokespeople is not uniform: a segment of consumers evaluates AI-generated personas on machine-appropriate criteria and responds more positively as a result.

Practitioner behavior reflects an industry recalibrating after early adoption. A Digiday analysis found that the share of marketing campaigns including AI or virtual influencers fell from 86% in 2024 to approximately 60% in 2025, a decline consistent with brands reassessing their use of the function after encountering consumer resistance. A Meltwater analysis of social media conversation in the first five months of 2025 found that highly engaged criticism was prominent in specific spikes of conversation, even though the broader social conversation remained polarized rather than uniformly negative. The function continues to attract investment, but the early enthusiasm among practitioners appears to have been tempered by experience of how consumers actually respond.

Taken together, the literature on virtual brand spokespeople describes a function where average consumer sentiment is often negative, but where responses vary substantially with disclosure, message context, perceived authenticity, source credibility, and individual differences. Rather than being uniformly rejected, virtual spokespeople appear to face more specific and demanding conditions for consumer acceptance than the other three functions examined in this research.

5. Regulatory and Industry Context

The four functions examined in this research do not operate in a regulatory vacuum. Disclosure requirements are tightening around AI-generated and AI-assisted content, and the compliance horizon is close enough that practitioners making deployment decisions today will need to account for it. Two developments are directly relevant.

The EU AI Act, which entered into force in August 2024, establishes a tiered framework for regulating AI systems based on risk level. For consumer-facing marketing, the most relevant provisions concern transparency obligations: when consumers interact with AI systems, they must be informed that they are doing so, unless the AI is obvious from context. Requirements for disclosing AI-generated or AI-manipulated content, particularly deepfakes, apply from August 2, 2026. For practitioners operating in or marketing to European consumers, this means that functions currently deployed without explicit AI disclosure, including virtual spokespeople and AI-generated advertising content, will face mandatory disclosure requirements within the compliance window. The Act mandates consumer awareness of AI deployment.

In January 2026, the Interactive Advertising Bureau released its AI Transparency and Disclosure Framework, the first industry-level attempt to standardize disclosure practices for AI use in advertising in the U.S. market. The framework takes a risk-based approach, calibrating disclosure requirements to the degree of AI involvement and the potential for consumer harm or confusion. It explicitly references the disclosure preferences documented in IAB's own consumer sentiment research. IAB's accompanying 2026 research argues that clearer communication and more consistent

disclosure practices can improve consumer attitudes toward AI-generated advertising and purchase likelihood, a finding that complicates the assumption that disclosure necessarily reduces commercial effectiveness. Consistent with the IAB framework's emphasis on consumer trust and regulatory risk, proactive disclosure may also confer a trust-signaling benefit, echoing the smaller positive pathway identified by Koning and Voorveld (2025).

For the four functions in this research, the regulatory picture is not uniform. AI customer service chatbots and AI personalized content delivery are already subject to general data protection and algorithmic transparency requirements in multiple jurisdictions. The new disclosure requirements bear most directly on AI creative advertising and AI virtual spokespeople, where AI involvement is less obvious to consumers and where the commercial intent is most salient. These are also, as the literature reviewed in Section 4 documents, the functions where consumer resistance is most concentrated and where the credibility cost of AI disclosure is highest. Practitioners deploying these functions in the coming year will be doing so inside a tightening regulatory frame that makes the question of how to disclose, not just whether to disclose, an operational priority. The consumer survey in this research uses a full-disclosure design across all four scenarios, making the findings directly relevant to a regulatory environment where disclosure is increasingly mandatory rather than optional.

6. Methodology

This white paper draws on two complementary research streams: a scenario-based consumer survey and practitioner evidence collected through a formal interview, field observation, and solicited industry responses. The two streams were designed independently. Their findings are interpreted together in Section 9 and translated into deployment recommendations in Section 10.

6.1 Consumer Survey Design

The core methodological choice was to put all four AI functions in front of the same consumer. Rather than running four separate studies with four separate samples, this research used a within-subjects design: each of the 150 participants read and evaluated all four scenarios. That design enables direct comparison. If participant A rates the chatbot higher than the virtual spokesperson, that difference is less likely to be an artifact of who happened to be assigned to which condition. No existing single-function study can produce that comparison. This one can.

Participants were recruited through Prolific, an online research platform, and data collection took place between March 13 and 14, 2026. To be eligible, participants had to be 18 or older, currently living in the United States, and occasional online clothing shoppers. A fourth screen asked whether they recognized the brand Arvo. Anyone who said yes was excluded. Arvo is a fictional brand created for this study to eliminate the influence of pre-existing brand associations on participant responses.

Each participant read four scenarios in randomized order. Randomization prevents any single scenario from benefiting or suffering simply because it appeared first or last.

Each scenario described the same fictional brand using AI in a different marketing function: generating a creative advertisement, personalizing the app homepage, powering a customer service chatbot, and deploying a virtual brand spokesperson named Kael. The scenario format and length were held constant across all four, with each running approximately 61 to 93 words. All four scenarios featured the same fictional apparel brand, the same product category, and the same first-person browsing or shopping context. AI disclosure was embedded in each scenario using consistent phrasing: the AI origin of each function was stated explicitly and directly, with no variation in disclosure depth or tone across conditions. Only the AI application changed.

After reading each scenario, participants rated the brand on three outcome measures. Brand attitude was measured with a five-item, seven-point semantic differential scale adapted from Spears and Singh (2004). Participants rated the brand on five pairs of opposing adjectives: Bad/Good, Unappealing/Appealing, Unpleasant/Pleasant, Unfavorable/Favorable, and Unlikeable/Likeable. Perceived authenticity was measured with a three-item, seven-point Likert scale adapted from Shoenberger and Kim (2023) and based on Beverland and Farrelly (2010). The three items were: “Arvo’s use of AI in this context seems truthful,” “Arvo’s use of AI in this context seems authentic,” and “Arvo’s use of AI in this context seems real.” Purchase intention was measured with a five-item, seven-point Likert scale adapted from Spears and Singh (2004). The five items were: “I intend to purchase Arvo,” “I predict I would buy Arvo,” “I would consider buying Arvo,” “It is likely that I will buy Arvo,” and “I am willing to buy Arvo.” All scales used endpoint labels only, and composite scores were computed by averaging items within each scale.

A total of 175 participants began the survey. Twenty-five were screened out for claiming familiarity with Arvo. The final sample of 150 completed the survey in full. A manipulation check question followed each scenario to confirm that participants correctly identified the AI function described. All 150 passed all four checks. Two attention check items were embedded across the survey. All 150 passed both. No further responses were excluded from analysis.

Scale reliability was assessed using Cronbach's alpha, the standard measure of internal consistency for multi-item scales (Cronbach, 1951; Shavelson & Steedle, 2022). A threshold of .70 is the conventional minimum for acceptable reliability. All twelve scales in this study ranged from .938 to .990, well above that threshold.

The primary analysis used one-way repeated measures ANOVA to test whether consumer sentiment differed significantly across the four AI functions. Before reporting those results, two standard assumption checks were conducted. Normality was assessed using the Shapiro-Wilk test. Sphericity was assessed using Mauchly's test. Where sphericity was violated, degrees of freedom were corrected using the Greenhouse-Geisser method. Post-hoc pairwise comparisons used Bonferroni correction to control for the increased probability of false positives when running multiple comparisons. Effect sizes are reported as partial eta-squared, with .01 representing a small effect, .06 a medium effect, and .14 a large effect (Cohen, 1988).

After completing all four scenarios, participants had the option to answer one open-ended question: "Out of the four scenarios you just read, which one felt most natural to you, and which felt most uncomfortable? Feel free to explain why." The question was optional. One hundred and forty-two of 150 participants responded. Because scenario

order was randomized, any references to ordinal position in participant responses were disregarded. Analysis focused entirely on semantic content. Responses were analyzed using reflexive thematic analysis (Braun & Clarke, 2006), with themes developed inductively from the data rather than applied from a pre-existing framework. Coding proceeded in two stages. In the first stage, a keyword-driven frequency count identified eight themes across all 142 responses, covering familiarity, uncanny valley reactions, job displacement concern, authenticity and deception concern, utility, privacy and data tracking concern, human fallback preference, and disclosure awareness. In the second stage, three sub-themes requiring finer semantic judgment, product representation concern in AI creative advertising, clothing representation concern in the virtual spokesperson condition, and data tracking discomfort in personalized content, were verified by reading each flagged response in full to confirm the coding matched the respondent's actual meaning. Quotes selected for presentation in Section 7 were chosen to represent the range of language used within each theme, not to amplify the most extreme responses. Thematic findings are presented in Section 7 alongside the quantitative results and referenced in Section 9 where they help explain the patterns the numbers alone do not fully account for.

The survey closed with a demographics section covering age, gender, race and ethnicity, education, household income, social media use, familiarity with AI in marketing, prior chatbot experience, and prior exposure to virtual influencers. Participants also rated their overall attitude toward AI in marketing on a seven-point scale. These variables are reported descriptively in Section 7 and used in

supplementary analyses to examine whether background characteristics moderated the main findings.

6.2 Practitioner Interview Design

The second research stream shifts the lens from consumer sentiment to practitioner behavior. Practitioner evidence was collected from three complementary sources, each contributing a distinct perspective. Each source contributes a distinct perspective and carries distinct limitations that shape how the findings are interpreted.

The primary source is one semi-structured interview with Steve Fowler, Head of Games at Ayzenberg Accelerate and Adjunct Professor at USC Annenberg. Fowler brings agency-side experience working with major gaming clients including Riot Games, WB Games, and Crystal Dynamics. The interview was conducted via video call, recorded with consent, and transcribed in full. It covered four areas: how the practitioner currently deploys AI across consumer-facing marketing functions, what factors shape deployment decisions, whether consumer sentiment enters that calculus, and how core versus mass-market audience segments respond differently to AI use in gaming marketing contexts.

The second source is a non-participant field observation conducted at a mid-sized mobile gaming publisher's marketing team event in April 2026. The researcher attended as a student observer and posed one open question during a Q&A segment.

Participants were not informed that responses would be used for research purposes.

For this reason, no individuals or organizations are identified in the findings, and observations are reported only at the level of general industry pattern. This material is treated as background context rather than primary evidence.

The third source is solicited practitioner commentary collected via two posts on r/AskMarketing and r/DigitalMarketing between March and April 2026. The researcher posted structured questions about AI deployment practices and consumer backlash experiences. Four responses were received from self-identified marketing practitioners. Respondents are anonymous, and their professional identities cannot be verified. This material is treated as secondary observational evidence, useful for triangulation but not for independent claims.

Analysis across all three sources focused on thematic synthesis. The goal was to identify patterns that reinforce or complicate the consumer survey findings. All practitioner findings are therefore reported as illustrative and exploratory.

7. Findings: Consumer Survey

One hundred and fifty U.S. adults completed the survey. The sample skewed toward younger adults, with 67% between the ages of 25 and 44. 59% identified as male and 41% as female. 45% held a bachelor's degree as their highest level of education, and 15% held a graduate or professional degree. Household income was distributed broadly, with the largest single group (29%) reporting between \$50,000 and \$74,999. The median survey completion time was eight minutes.

Participants were generally familiar with AI in marketing contexts. 49% described themselves as very or extremely familiar with AI being used in marketing or advertising. 95% had previously interacted with an AI chatbot for customer service. 55% had seen a virtual influencer on social media before. Overall attitude toward AI in marketing, rated on a scale from 1 (very negative) to 7 (very positive), averaged 3.46 (SD = 1.74), sitting slightly below the neutral midpoint. The sample, in other words, was experienced with AI but not particularly enthusiastic about it.

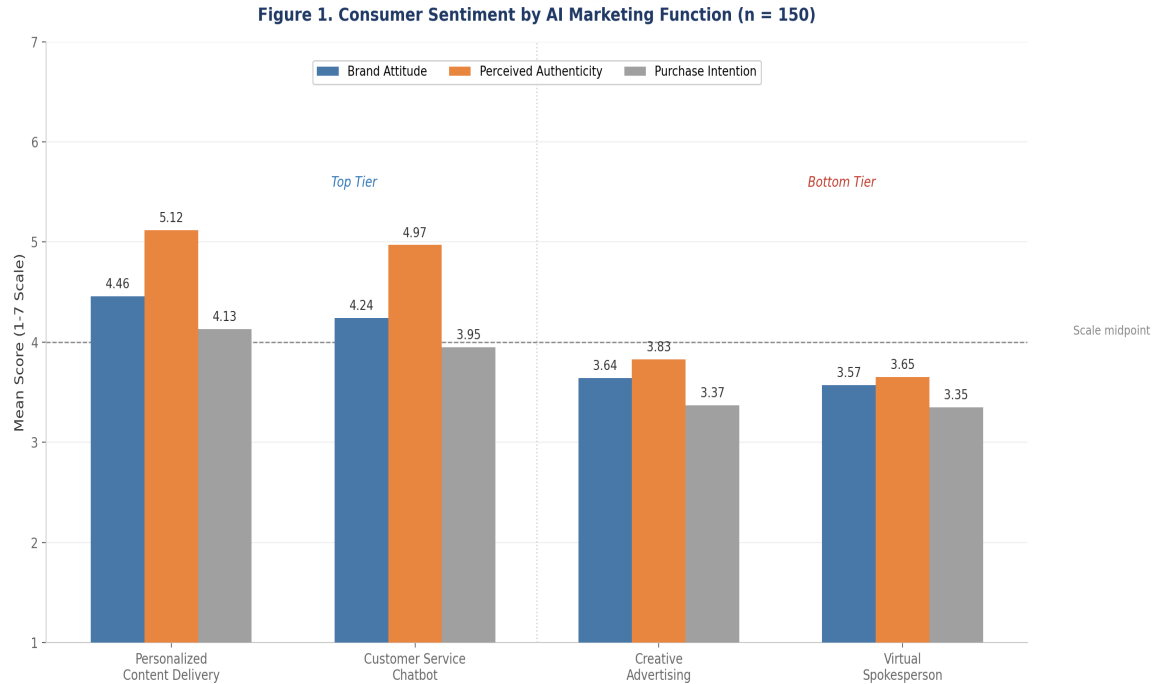
7.1 Cross-Function Results

Table 1 shows the mean scores and standard deviations for all three outcome measures across the four AI marketing functions. All scales ran from 1 to 7, with higher scores indicating more positive sentiment. Full scale items are described in Section 6.1.

AI Marketing Function	Brand Attitude M (SD)	Perceived Authenticity M (SD)	Purchase Intention M (SD)
AI Creative Advertising	3.64 (1.65)	3.83 (1.78)	3.37 (1.58)
AI Personalized Content Delivery	4.46 (1.51)	5.12 (1.38)	4.13 (1.60)
AI Customer Service Chatbot	4.24 (1.49)	4.97 (1.51)	3.95 (1.61)
AI Virtual Spokesperson	3.57 (1.77)	3.65 (1.85)	3.35 (1.70)

All scales 1-7. Higher scores indicate more positive sentiment. n = 150.

Table 1. Consumer sentiment by AI marketing function (n = 150)



A one-way repeated measures ANOVA confirmed that differences across the four functions were statistically significant for all three outcome measures: brand attitude ($F(2.45, 365.74) = 36.50, p < .001, \text{partial eta-squared} = .197$), perceived authenticity ($F(2.45, 364.39) = 70.77, p < .001, \text{partial eta-squared} = .322$), and purchase intention ($F(2.70, 402.13) = 35.39, p < .001, \text{partial eta-squared} = .192$). All three effect sizes fall in the large range. The AI function a brand chooses accounts for a meaningful share of how consumers feel about it.

7.2 What the Pattern Shows

Post-hoc comparisons reveal a consistent two-tier structure across all three outcome measures. AI personalized content and AI customer service chatbots score significantly higher than AI creative advertising and AI virtual spokespersons. The gap between the two tiers is significant across all three measures ($p < .001$ in each case). Within each tier, differences are small and statistically non-significant: personalized content and

customer service are not meaningfully different from each other, and neither are creative advertising and virtual spokesperson. The three outcome measures are also highly correlated within each function, with Pearson r values ranging from 0.65 to 0.90, all within the large range by Cohen's (1988) conventions. This suggests that the three scales are moving together across functions, and that the pattern reported here holds across multiple evaluative dimensions rather than being specific to any one scale.

Perceived authenticity shows the sharpest separation between tiers. Personalized content scored 5.12 on perceived authenticity and virtual spokesperson scored 3.65, a gap of 1.47 points on a seven-point scale. Brand attitude and purchase intention show similar patterns but slightly smaller gaps. Of the three outcome measures, perceived authenticity is the most sensitive to which AI function is in use, a pattern that Section 9 examines in detail.

7.3 Function-by-Function Summary

AI Personalized Content Delivery

Personalized content ranked first across all three measures. Brand attitude averaged 4.46, perceived authenticity 5.12, and purchase intention 4.13. It is the only function where perceived authenticity exceeded 5.0, crossing the threshold between neutral and positive.

AI Customer Service Chatbot

The customer service chatbot ranked second on all three measures, close behind personalized content. Brand attitude averaged 4.24, perceived authenticity 4.97, and purchase intention 3.95. The gap between personalized content and customer service is

small and does not reach statistical significance on any measure, meaning these two functions are effectively tied at the top. Where the chatbot diverges from personalized content in the open-ended responses is in the quality of that acceptance: chatbot acceptance is more conditional, grounded in familiarity with a widely encountered tool rather than positive evaluation of the experience itself.

AI Creative Advertising

AI creative advertising ranked third. Brand attitude averaged 3.64, perceived authenticity 3.83, and purchase intention 3.37. All three scores sit below the scale midpoint of 4.0, marking a clear shift from the top two functions. Creative advertising scored significantly lower than both personalized content and customer service on all three measures ($p < .001$), and was statistically indistinguishable from the virtual spokesperson on all three.

AI Virtual Spokesperson

The virtual spokesperson ranked last. Brand attitude averaged 3.57, perceived authenticity 3.65, and purchase intention 3.35. These are the lowest scores across the board, and also the most dispersed: standard deviations of 1.77, 1.85, and 1.70 respectively indicate that consumer reaction to the virtual spokesperson was more polarized than to any other function. Some participants responded positively. More responded negatively. The virtual spokesperson also produced the strongest negative language in open-ended responses, discussed in Section 7.4.

7.4 Supplementary Consumer Insights

Familiarity as the basis for acceptance. 23% of respondents described a function as natural primarily because they had encountered it before. The majority of those

references pointed to the chatbot, where acceptance tended toward resignation rather than endorsement.

“It’s not a great experience, but one I’ve come to accept as normal.”

Utility as a differentiator for personalized content. 18% of respondents described a function positively on the basis of practical helpfulness, and the large majority of those references pointed to personalized content. Where chatbot acceptance was passive, personalized content acceptance was functional. Participants described it as genuinely helpful rather than merely tolerable.

“The one that felt the most natural was helping me choose products based on past purchases. This is something that seems quite useful.”

The virtual spokesperson as a unique category of discomfort. 27% of respondents used language associated with the uncanny valley effect when describing the virtual spokesperson, including words like “creepy,” “eerie,” “dystopian,” and “uncanny valley” itself. No other function triggered this language. 21% expressed concern about authenticity or deception more broadly, but those references concentrated heavily on the virtual spokesperson. Several connected their discomfort explicitly to the deceptive potential of a lifelike AI character presented as a brand representative.

“The AI model used was portrayed as a ‘human’ on their social media page. It is not a real person, but it is being advertised as such.”

Job displacement as a recurring concern. 14% of open-ended respondents mentioned job displacement unprompted, primarily in response to the virtual spokesperson and AI creative advertising. This concern did not appear in the quantitative scales, which measured brand-level sentiment rather than broader social attitudes. Its presence in the open-ended data suggests that some of the negative scores for these two functions carry an ethical dimension that the rating scales do not fully surface.

“It not only takes jobs away from models and photographers, but it probably depicts unrealistic beauty standards.”

8. Findings: Practitioner Evidence

This section draws on three sources: a semi-structured interview with Steve Fowler, Head of Games at Ayzenberg Accelerate; field observation at an anonymized mobile gaming publisher; and solicited practitioner commentary collected via Reddit. Each source is described where it is used.

8.1 AI Creative Advertising

Practitioner evidence on AI creative advertising points to a consistent gap between back-end AI use and front-end disclosure. Across all three sources, AI is heavily integrated into creative workflows. What reaches the consumer is a different matter.

Fowler described a clear internal norm among his gaming clients: AI is acceptable for concepting, not for final consumer-facing output. In practice, his team uses generative tools to produce up to 15 layout variations of a campaign concept within hours, presents them to the client for directional feedback, and then rebuilds the approved direction using licensed brand assets. The final piece shown to the public contains no AI-generated imagery. The stated reason is IP liability. Major publishers cannot risk litigation over training data provenance, particularly when their character art is legally protected. Consumer sentiment does not appear in this calculation.

Field observation at the anonymized publisher suggested a similar pattern. Practitioners described using AI to generate lighting for trailers and PR materials, to visualize brand collaboration pitch decks, and to draft copy that a human editor then refines. In each case, the AI output is an intermediate product. The consumer-facing result is human-reviewed. No one mentioned proactive disclosure to consumers as part of this workflow.

One anonymous Reddit respondent offered a telling data point. In a campaign that was approximately 80% AI-assisted, consumers praised what they called the “authentic human touch” of the work (r/AskMarketing, 2026). The respondent described this as evidence that consumers generally cannot identify AI creative unless the execution is conspicuously poor. The same respondent characterized the current industry norm as “don’t ask, don’t tell”: AI use is widespread but rarely surfaced unless a journalist or a vocal consumer forces the question.

The risks of consumer-facing AI creative are illustrated by Coca-Cola’s 2024 Christmas campaign, in which the company released a fully AI-generated remake of its classic “Holidays Are Coming” commercial. The ad drew immediate backlash, with viewers describing it as “soulless” and “devoid of creativity” (Horvath, 2024). Coca-Cola declined to retract the campaign, stating it remained committed to exploring the intersection of human creativity and technology. The episode is notable because Coca-Cola explicitly positioned AI as the primary creative force rather than a background production tool, and the backlash followed directly from that positioning.

The Crystal Dynamics case offers a contrast. When the studio disclosed on Steam that AI-assisted tools had been used to upscale textures in its *Legacy of Kain* remaster, a vocal minority of core gamers reacted negatively. According to Fowler, a larger segment of the fanbase pushed back against the criticism, arguing that using AI to upscale existing human-authored artwork was categorically different from using AI to generate new content (S. Fowler, 2026). The disclosure was mandatory under Steam’s platform policy. Steam user reviews for the title reached 97% positive at launch, suggesting the AI disclosure did not meaningfully affect overall player reception (Crystal Dynamics,

2026). The contrast with Coca-Cola suggests that scope matters: AI used as a production efficiency tool within a human-authored creative work generates a different response than AI positioned as the primary author.

8.2 AI Customer Service Chatbots

Chatbots generated the most consistent practitioner acceptance across all sources.

Fowler described Discord moderation bots as routine in gaming communities, widely understood, and rarely contested. Their function is factual: they enforce community rules, answer documented questions, and route users to the right channels.

Field observation at the anonymized publisher identified a deployed AI agent handling what one participant described as Level 0 player support queries: password resets, account access, and basic troubleshooting. Complex issues are escalated to human agents. No participant described consumer resistance to this system.

The Reddit evidence introduces a more complicated picture. One respondent who moderates a financial company's online community described what happened when the company replaced a significant portion of its customer support team with an AI system (r/AskMarketing, 2026). The community was flooded with unresolved support requests. The moderator closed the subreddit temporarily. When the company proposed deploying AI agents to respond within the community itself, the moderator banned them outright. The company eventually acknowledged the transition had failed. The respondent was explicit about the root cause: the backlash was not directed at the AI. It was directed at the layoffs that preceded it. Consumers were angry because humans had been removed, and the AI became a visible symbol of that decision rather than the source of grievance in itself.

This distinction between AI as a tool and AI as a labor replacement signal appears to be the critical variable. When a chatbot supplements human capacity, consumer response is generally neutral. When it is deployed in place of humans who were visibly let go, the chatbot inherits the reputational cost of that decision.

8.3 AI Virtual Brand Spokespeople

Practitioner evidence on virtual brand spokespeople centers on two publicly documented cases from the gaming sector, both referenced by Fowler during the interview. Neither maps precisely onto the category as defined in Section 4: one involves AI-generated game commentary, the other an LLM-driven IP character. Both are included here because each deploys an AI persona to interact with or perform for consumers in ways that parallel the spokesperson function.

The first is *The Finals*, a multiplayer shooter released in late 2023 by Embark Studios. The game's real-time commentators and contestant voice lines were generated using AI text-to-speech rather than human voice actors. When the studio confirmed this in a podcast, the decision drew criticism from both players and professional voice actors. Players criticized the AI voices as unconvincing (Randall, 2023), while professional voice actors objected to the displacement of human performers from paid work (Middler, 2023). Embark defended the choice on efficiency grounds, clarifying that the AI voices were modeled on recordings from real actors rather than generated from scratch (Randall, 2023). Despite the controversy, the game attracted over 260,000 concurrent players during its opening weekend, suggesting the backlash did not deter the broader playerbase (Ingram, 2023).

The second case introduces a different category of risk. In May 2025, Fortnite integrated an AI-powered Darth Vader character driven by Google's Gemini 2.0 Flash model, voiced using a digital replica of the late James Earl Jones. Players could interact with the character in real time using voice. Shortly after launch, users had manipulated the model into producing racial slurs and homophobic language (PinkNews, 2025). Epic Games deployed hotfixes rapidly, but the content had already circulated widely online. Three days after launch, SAG-AFTRA filed an unfair labor practice charge, alleging that Epic had replaced union voice actors without notice or negotiation (Game Informer, 2025). The case surfaces two distinct risks simultaneously: an alignment risk, where an LLM persona operating in an open consumer environment cannot reliably stay within intended behavioral bounds, and a labor displacement risk, where using AI to replicate a human performer's voice generates institutional and legal backlash independent of consumer response.

Together, these cases suggest that AI virtual spokespeople face a compound challenge. The Finals shows that even a commercially successful deployment can attract industry-level criticism around labor displacement, even when consumer uptake remains strong. The Fortnite case adds a second layer: when an LLM-driven persona operates in an open consumer environment, behavioral alignment cannot be guaranteed, and the consequences of failure are immediate and public. Brands considering this function must account for both.

8.4 AI Personalized Content Delivery

This function generated the least discussion across all three practitioner sources. Field observation at the anonymized publisher surfaced AI-assisted multilingual localization

as one operational use, though this is more accurately described as a production tool than a personalization system. Reddit commentary did not address this function directly.

The relative absence of practitioner evidence here is itself informative. Personalized content delivery in the gaming sector operates primarily at the platform level. Steam's recommendation engine, for instance, is driven entirely by player behavior data rather than inputs from individual publishers, and Valve has explicitly stated that developers cannot meaningfully influence how the algorithm surfaces their games to players (McAloon, 2019). Individual brands and agencies interact with these systems as participants, not architects. This may explain why none of the practitioner sources framed personalized content delivery as a function they actively deploy, in contrast to the other three functions.

8.5 Cross-Function Patterns

Three patterns emerge across the four functions.

The first is that IP liability, not consumer sentiment, is the primary constraint on AI creative advertising adoption among the practitioners represented here. Brands and agencies want to use generative AI for final creative output but cannot take on the litigation risk of unverifiable training data provenance. In most cases, consumers are not even aware that AI was involved in the creative process.

The second is that non-disclosure is the current industry default. Across all sources, proactive AI disclosure to consumers is the exception rather than the rule. Practitioners disclose when legally required to, as under Steam's platform policy, or when a backlash forces the question. The Reddit evidence suggests that even heavily AI-assisted

campaigns can pass as human-made when execution quality is high (r/AskMarketing, 2026).

The third is that the actual trigger for consumer backlash is often labor displacement rather than AI use per se. The financial community case and the SAG-AFTRA response to Fortnite both illustrate that consumers and industry stakeholders tolerate AI as a production tool but react when it is framed as, or deployed as, a replacement for identifiable human workers. This pattern is consistent with the consumer survey finding that functions perceived as replacing a human role generate stronger negative reactions than functions perceived as tools.

9. Discussion

9.1 Consumer Survey

The survey data show a two-tier structure across all three outcome measures.

Personalized content delivery and customer service chatbots score consistently higher than the other two functions, with most measures at or above the scale midpoint. AI creative advertising and AI virtual spokespeople score below the midpoint across all three measures. All cross-tier gaps are statistically significant at $p < .001$. Within each tier, differences are not. The same split appears on brand attitude, perceived authenticity, and purchase intention. Three features of this pattern require explanation: why authenticity scores diverge so sharply despite identical disclosure conditions across all four scenarios, why the two top-tier functions earn acceptance in qualitatively different ways despite statistically similar scores, and why the two bottom-tier functions score equally low through distinct mechanisms. The theoretical frameworks built in Section 3 each contribute part of the explanation.

Ruling Out Familiarity

Before turning to theory, one alternative explanation needs to be ruled out. The sample was highly experienced with AI: 95% had previously used a customer service chatbot, and 49% described themselves as very or extremely familiar with AI being used in marketing. These two figures measure different things. The first is a behavioral indicator specific to one function. The second is a self-reported general familiarity that applies across all four. Taken together, they describe a sample that was not encountering AI marketing as a novelty. A reasonable hypothesis follows: consumers who are more familiar with AI may simply respond more positively across the board, which would

produce the two-tier pattern as a byproduct of experience rather than anything structural about the functions themselves.

To test this, participants were divided into a high-familiarity group (those who identified as extremely or very familiar with AI in marketing, $n = 73$) and a low-familiarity group (those who identified as slightly familiar or not familiar at all, $n = 25$). Moderately familiar participants ($n = 52$) were excluded from this comparison to create a clean contrast between the two ends of the familiarity distribution. Independent samples t-tests comparing the two groups on all twelve outcome measures (brand attitude, perceived authenticity, and purchase intention for each of the four functions) returned no significant differences, with all $ps > .05$ and no correction applied given the exploratory nature of the comparison. For the virtual spokesperson specifically, a parallel test using prior virtual influencer exposure as the grouping variable reached the same conclusion: participants who had previously seen virtual influencers on social media ($n = 82$) rated the virtual spokesperson no higher than those who had not ($n = 40$), with no significant differences on brand attitude, perceived authenticity, or purchase intention. Across both measures and all four functions, the two-tier structure holds regardless of what prior experience a participant brings. One caveat applies to personalized content specifically: the survey did not include a behavioral measure of prior personalized recommendation use, so the familiarity argument for that function rests on the general AI marketing familiarity score rather than a direct indicator. This is noted as a limitation in Section 11. The broader conclusion holds: familiarity does not account for the two-tier pattern. Something about the functions themselves does.

Why Authenticity Scores Diverge Despite Identical Disclosure

The three outcome measures are highly correlated within each function, with within-function Pearson r values ranging from 0.65 to 0.90, all within the large range by Cohen's (1988) conventions for Pearson's r . This means the cross-function effect is consistent across brand attitude, perceived authenticity, and purchase intention, and is not an artifact of any single scale. Of the three, perceived authenticity shows the largest effect size (partial eta-squared = .322), and the widest cross-tier gap: personalized content scored 5.12 and the virtual spokesperson scored 3.65, a 1.47-point difference on a seven-point scale. All four scenarios disclosed AI involvement in identical terms, yet authenticity scores diverge sharply. Disclosure alone cannot account for this pattern. The category of source each function places the AI in offers a more complete explanation.

Sundar and Nass (2001) distinguish between technological sources, systems that consumers perceive as functional and non-human, and visible sources, agents presented with human characteristics and social roles. The distinction matters for a specific reason: consumers bring different evaluative standards to each category. A technological source tends to be assessed on whether it works. A visible source tends to be assessed on whether it can be trusted as a person.

Chatbots and personalized content operate as technological sources. Consumers know they are interacting with a system. The evaluative frame is functional: did it work, was it helpful? Authenticity in the personal-identity sense is simply not a criterion consumers apply to a recommendation engine or a service bot, so disclosure changes nothing on that dimension. Virtual spokespeople occupy a fundamentally different position. Kael is presented as a brand representative with a name, a visual identity, and a social role

previously held by human endorsers. This places Kael in the visible source category, where the evaluative criteria shift entirely. Koles, Audrezet, Moulard, Ameen, and McKenna (2024) find that consumers evaluate virtual influencers on three authenticity dimensions: whether the persona is transparent and non-deceptive about its artificial nature (true-to-fact authenticity), whether its attributes correspond with a socially recognized ideal or role standard for what such an entity should be (true-to-ideal authenticity), and whether the entity behind it acts from genuine intrinsic motivation rather than pure commercial interest (true-to-self authenticity). For Kael, the second and third dimensions present structural difficulties. As a persona created and controlled by a brand, Kael's role is defined by commercial objectives, and the question of whether it matches any independent social ideal for a spokesperson is complicated by the fact that there is no established human standard a brand-created AI character naturally corresponds to. Consumers who apply these authenticity criteria, as 21% did explicitly in the open-ended responses, find the function difficult to accept not because the AI fails to perform but because the role it occupies demands a kind of authenticity that a brand-created character cannot credibly claim. Perceived authenticity for the virtual spokesperson ($M = 3.65$) reflects this evaluative mismatch. Creative advertising sits between the two: the AI is not a visible entity. It originates content that carries markers of human creative judgment. Grigsby, Michelsen, and Zamudio (2025) found that the credibility cost of AI disclosure concentrates in ad elements consumers use to infer human judgment and warmth, with trust and ad attitudes recovering when AI is used only for tangible elements while intangible elements remain human-produced. In Grigsby et al.'s framework, intangible elements are those that communicate relational

qualities, human taste, and creative judgment rather than objective physical attributes. For a clothing ad, styling choices and the way products are presented may serve this function: consumers can use them to infer that a human made considered aesthetic decisions about what looks good and why. When those elements are disclosed as AI-generated, that inference may weaken, consistent with Grigsby et al.'s finding that disclosure is more damaging when it affects elements consumers use to infer human judgment and warmth. This may help explain the authenticity cost visible in this study's creative advertising scores.

The findings are consistent with source attribution theory: where consumers never expected a human source, the authenticity cost of disclosure is substantially reduced. Where they encounter something in a human source role, the same disclosure may activate a set of criteria that an AI-generated entity may find difficult to satisfy.

Why the Top Tier Earns Acceptance in Different Ways

Brand attitude for personalized content ($M = 4.46$) and chatbot ($M = 4.24$) are statistically indistinguishable, as are purchase intention scores (4.13 vs 3.95). The quantitative scores treat them as equivalent. The open-ended responses show otherwise.

Twenty-three percent of respondents described a function as natural primarily because they had encountered it before, and the majority of those familiarity references pointed to the chatbot. The acceptance is passive: participants described it as “not a great experience, but one I’ve come to accept as normal” and noted that chatbots are standard on most websites “whether they are helpful is another story.” The function is

tolerated because it is familiar, not because it is valued. Personalized content attracted qualitatively different language. Eighteen percent of respondents cited a function as natural on the basis of practical helpfulness, and the large majority of those utility references pointed to personalized content. Participants described it as “quite useful” and “helpful” and framed the AI as working for them rather than merely being present.

Construal Level Theory (Trope & Liberman, 2010) explains why both functions earn acceptance at all. Both involve psychologically proximate interactions: the chatbot is real-time and personally addressed, and the recommendation engine draws on the consumer’s own recent behavior. CLT offers a consistent account: psychological proximity tends to activate concrete, feasibility-based evaluation, making functional concerns more salient than questions about whether the AI is the right kind of agent for the job. In many routine cases, a well-functioning AI can address those functional concerns as effectively as a human.

The Capability-Personalization Framework (Qin et al., 2025) offers an account of why the two functions earn that acceptance for different reasons. The chatbot in a routine service context fits the appreciation condition cleanly: AI is perceived as capable, and human-like personalization is unnecessary. Personalized content is more complicated. By definition, it is a context where personalization is central, which the framework would predict should increase aversion risk. Yet it ranks first across all three measures. The resolution, consistent with Hardcastle, Vorster, and Brown (2025), lies in how consumers frame the interaction. When AI uses past purchases to surface relevant options, most consumers evaluate it on utility grounds rather than autonomy grounds: the AI is perceived as working for them, reflecting their own stated preferences back.

That framing helps explain why the qualitative language around personalized content appears more enthusiastic than the language around chatbots. The framing is partly under practitioner control, and it matters.

Why the Bottom Tier Scores Low Through Different Routes

Creative advertising (brand attitude $M = 3.64$, perceived authenticity $M = 3.83$, purchase intention $M = 3.37$) and virtual spokesperson (3.57, 3.65, 3.35) are statistically indistinguishable across all three measures. Both fall below the scale midpoint. The quantitative scores treat them as equivalent failures. The open-ended responses reveal that they arrive at the same outcome through distinct mechanisms.

What both functions share is the combination of AI disclosure with transparent commercial intent. The Persuasion Knowledge Model (Friestad & Wright, 1994), extended to AI contexts by Watson, Valsesia, and Segal (2024), specifies that AI-related persuasion knowledge activates when two conditions are simultaneously present: the consumer recognizes the agent as AI, and perceives that it carries persuasive intent. In this study, all four scenarios disclosed AI explicitly, satisfying the first condition across the board. The second condition is where the functions diverge. For chatbots and personalized content, the function is primarily instrumental. For creative advertising and virtual spokespeople, the commercial intent is the point. Both preconditions are most clearly and most strongly met for the bottom tier. Accordingly, scrutiny is more likely to intensify there.

The scrutiny has a specific cost. Bui (2025) found that AI disclosure in advertising contexts reduced perceived credibility, which reduced both ad value and purchase intent. Koning and Voorveld (2025) found that disclosure reduced trust in both the

advertisement and the organization behind it. One interpretation of these disclosure effects is substitution: consumers may use the human judgment implicit in advertising as a proxy for trustworthiness, and disclosure removes that proxy without fully replacing it. For chatbots and personalized content, this substitution mechanism is likely weaker because consumers generally expect less human creative judgment in the first place. This pattern is consistent with the view that the same disclosure condition costs more in one tier than the other.

Beyond the shared PKM mechanism, each bottom-tier function has its own additional problem. For creative advertising, 6 open-ended respondents raised a concern the rating scales did not capture: they could not tell whether the product shown in an AI-generated ad was an accurate representation of what they would receive. Responses described AI-generated clothing imagery as “fake,” “staged,” and failing to show “what it would actually look like if I purchased it.” This is a product representation concern. As discussed in the authenticity section above, Grigsby et al. (2025) find that AI disclosure is most damaging in ad elements consumers use to infer human judgment and care, and styling choices in clothing advertising may serve that function. A separate but related concern appeared in 3 responses about the virtual spokesperson: participants questioned whether an AI-generated figure could accurately convey how clothing would actually look on a real body. Both concerns target the practical purchasing utility of the function rather than AI ethics in the abstract.

The virtual spokesperson carries additional weight. Open-ended responses cited it as most uncomfortable by 47 respondents (33%), and the language they used identifies not one mechanism but four operating simultaneously. Respondents described uncanny

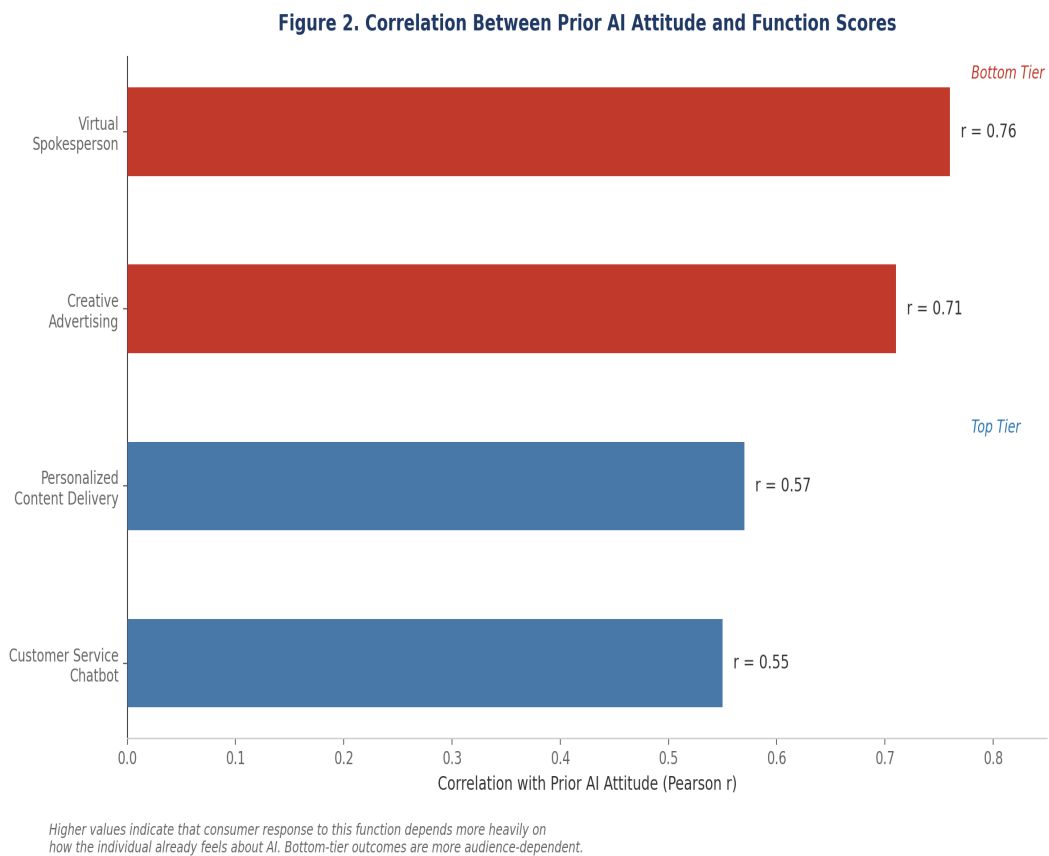
valley reactions to Kael as a lifelike character that “looks like a real person but isn’t real.” They raised job displacement unprompted, noting that AI models “take jobs away from models, photographers, etc.” They questioned brand authenticity, describing Kael as “not a real person, but being advertised as such.” And some connected their discomfort to product evaluation: an AI figure cannot wear clothes the way a real body does, raising doubt about what the clothing would actually look like. As discussed in the authenticity section above, Koles et al. (2024) find that consumers apply a differentiated authenticity framework to virtual influencers across three dimensions simultaneously. A fictional AI spokesperson encounters all three criteria at once, and four “distinct concerns compound on a single function. No single design adjustment resolves all of them.

As noted in the familiarity analysis above, participants who had previously seen virtual influencers on social media ($n = 82$) rated Kael no higher than those who had not ($n = 40$), with no significant differences on any outcome measure. In this sample, the resistance does not appear to be a simple novelty effect that normalizes with prior exposure.

Prior AI Attitudes Amplify Bottom-Tier Sensitivity

An analysis of the Pearson correlations between each participant’s overall attitude toward AI in marketing and their composite scores on each function shows a consistent asymmetry. For AI creative advertising, the average absolute correlation across the three outcome measures is 0.71. For the virtual spokesperson, it is 0.76. For personalized content and customer service chatbots, those figures are 0.57 and 0.55, respectively. By Cohen’s (1988) conventions for Pearson’s r , values at or above 0.50 fall in the large range. All four functions meet that threshold. The bottom-tier functions,

however, sit notably higher within that range (0.71 and 0.76) than the top-tier functions (0.57 and 0.55), indicating that individual AI attitude is a considerably stronger predictor of consumer response for creative advertising and virtual spokespersons than for chatbots or personalized content.



The top-tier functions are relatively robust to prior disposition. Most consumers extend reasonable acceptance to chatbots and personalized content regardless of how they feel about AI generally. The bottom-tier functions are more polarizing. Skeptical consumers become more skeptical when they encounter creative advertising or a virtual spokesperson. Favorable consumers respond more positively. Prior attitude amplifies the function-level signal in both directions. The commercial risk of bottom-tier deployment is therefore that consumer response is more strongly tied to prior AI

attitudes, making outcomes more dependent on audience composition than top-tier deployment.

A Conditional Reading of Personalized Content

A minority counter-current runs through personalized content. Three respondents expressed discomfort with the data collection underlying the feature, rather than with the AI itself: “It is like the company is tracking me.” Hardcastle et al. (2025) describe exactly this dynamic: when personalization is experienced as dataveillance rather than service, the same AI behavior that reads as helpful to one consumer reads as invasive to another. These three responses did not move the aggregate scores. But they signal that personalized content’s high acceptance is conditional. Consumers accept it when they read it as a tool working in their interest. When the data exchange behind it becomes salient, that reading can shift. How practitioners describe the feature and what they disclose about data use affects which reading consumers arrive at.

Familiarity with AI does not account for the two-tier pattern. What does is a set of function-level characteristics that shape which evaluative standards consumers bring to each interaction. Where AI operates as a background system, consumers evaluate it on whether it works. Where it occupies a role traditionally held by a person, authenticity criteria are activated, and identical disclosure conditions produce sharply different outcomes across the four functions. Psychological proximity shapes whether that evaluation is grounded in concrete feasibility or abstract value judgments, explaining why both top-tier functions earn acceptance while the quality of that acceptance differs. The visibility of commercial intent determines how readily persuasion knowledge activates, concentrating resistance in functions where AI is the visible author of a

persuasive message. And prior AI attitudes amplify all of these effects most strongly in the bottom tier, making audience composition a meaningful risk variable for practitioners deploying AI in visible, persuasion-oriented roles. The open-ended responses confirm that consumers can articulate these distinctions themselves.

9.2 Practitioner Evidence

Section 9.1 interpreted the consumer survey through the theoretical frameworks developed in Section 3. This section asks a different question: where does the practitioner evidence from Section 8 converge with, and diverge from, the consumer data? Four patterns emerge from this comparison.

Tool or Replacement: The Shared Fault Line

The consumer survey and the practitioner evidence arrive at the same dividing line through independent routes. In the survey, source attribution theory (Sundar & Nass, 2001) explains the two-tier split: consumers evaluate AI differently depending on whether they perceive it as a functional system or as something occupying a human role. The practitioner evidence identifies the same boundary from the deployment side. Across Section 8, the consistent trigger for consumer backlash is not AI use itself but the perception that AI is replacing a human performer.

Three cases from Section 8 converge on this point. The financial community case shows what happens when AI replaces human support staff and the service quality drops: the community shut down, the moderator banned the company, and every subsequent AI proposal was rejected before it could be evaluated on its own merits. The Finals shows that even a commercially successful AI deployment attracts reputational friction when the decision reads as displacing human voice actors. The

Fortnite Darth Vader case extends the pattern to an institutional level, with SAG-AFTRA filing an unfair labor practice charge over the replacement of union performers. In each case, the backlash targeted the displacement, not the technology.

These cases confirm what the consumer data predicts. The two-tier structure reflects a fundamental distinction in how consumers categorize AI: as a system performing a function, or as an entity occupying a role that was previously human. When deployment is coupled with visible labor displacement, even functions that consumers otherwise tolerate can inherit the reputational cost of the displacement. Functions that consumers already associate with human performers face this risk by default.

Audience Composition Shapes Bottom-Tier Outcomes

The consumer survey found that prior AI attitudes correlate more strongly with bottom-tier function scores ($r = 0.71$ for creative advertising, 0.76 for virtual spokesperson) than with top-tier scores (0.57 for personalized content, 0.55 for chatbot). Top-tier functions are relatively robust to audience disposition. Bottom-tier functions are more polarizing, with outcomes shaped by who is in the audience.

The practitioner evidence identifies a parallel pattern through a different mechanism. Fowler estimated that roughly 90 to 95% of mass market gamers are indifferent to whether AI was used in a game's production. The remaining 5 to 10%, core gamers with deep product investment and community ties, actively scrutinize AI deployment. This is not a general AI attitude. It is a product engagement variable. Core gamers care because they are invested in the medium, its creative workforce, and its production standards.

The two mechanisms are distinct, but their implications converge. The survey measures a general disposition toward AI in marketing. Fowler's observation captures a domain-specific identity tied to community membership and creative investment. Both point to the same practical conclusion: consumer response to bottom-tier functions is more sensitive to individual disposition, making audience composition a more consequential variable for deployment outcomes.

For practitioners, this means that average consumer sentiment scores understate the deployment risk of bottom-tier functions. The same creative AI campaign can succeed with one audience and fail with another. Average scores from a general sample cannot predict the outcome of a specific deployment, because audience composition determines how strongly prior attitudes shape the response.

The Personalized Content Blind Spot

Personalized content delivery scored highest across all three consumer measures: brand attitude ($M = 4.46$), perceived authenticity ($M = 5.12$), and purchase intention ($M = 4.13$). It is the function consumers respond to most favorably. It is also the function that generated the least practitioner discussion across all sources in Section 8.

In the gaming sector, personalized content delivery operates primarily at the platform level. Steam's recommendation engine, Epic's personalized storefront, and mobile app store algorithms all use AI to surface relevant content to individual users. These are platform-executed functions. The brands and publishers whose products appear in those recommendations exercise little direct control over the recommendation logic itself.

The consumer survey measured responses to a brand-executed version of this function: an AI system that analyzed a customer's purchase history to generate personalized product suggestions on the brand's own website. Consumers responded favorably. But the practitioner evidence suggests that in practice, most personalized content delivery happens through intermediaries that brands do not control.

This creates a specific tension. The AI function that earns the strongest consumer acceptance is also the one where brands have the least direct agency. Section 9.1 showed that personalized content's high scores are conditional: consumers accept it when they read the AI as working in their interest, and a minority counter-current emerges when the data exchange behind it becomes salient. How the feature is framed matters. But if the feature is primarily platform-executed, brands have limited ability to control that framing.

The Concepting Norm: Right Outcome, Wrong Reason

Section 8.1 documented a consistent practitioner norm across multiple sources: AI is acceptable for concepting and iteration, not for final consumer-facing creative output. AI accelerates the ideation phase while human execution produces the deliverable that reaches the consumer. Practitioners adopted this workflow for legal and operational reasons, primarily because IP ownership of AI-generated content remains unsettled and clients want to own their final creative assets without ambiguity.

From the consumer side, this workflow happens to address exactly the right problem. Grigsby, Michelsen, and Zamudio (2025) found that AI disclosure damages ad credibility most when it affects the elements consumers use to infer human judgment and warmth, and that trust recovers when AI handles only tangible, mechanical

elements while intangible elements remain human-produced. The concepting norm preserves those human judgment signals in the final output by keeping AI in the back end. The Reddit evidence from Section 8.1 illustrates the result: in a campaign that was approximately 80% AI-assisted, consumers praised what they perceived as an “authentic human touch.”

The alignment is instructive but accidental. Practitioners adopted this workflow to manage legal risk and discovered that it also mitigated consumer resistance, without necessarily understanding why. The consumer data explain the mechanism: by keeping AI in a tool role during the production process rather than making it the visible author of the final output, the concepting norm avoids triggering the source attribution and persuasion knowledge mechanisms that penalize bottom-tier functions.

10. Recommendations

The findings in Chapters 7 through 9 translate into a small set of decisions practitioners face when deploying AI in consumer-facing functions. The consumer survey identified a two-tier structure: personalized content delivery and customer service chatbots earned significantly higher consumer acceptance, while creative advertising and virtual spokespeople scored below the scale midpoint across all three measures. This chapter organizes deployment decisions in two layers: three cross-cutting principles that apply across all four functions, followed by function-specific checklists that turn the principles into concrete deployment judgments.

Cross-Cutting Principles

Treat replacement framing as the primary risk variable. Before deploying any AI function, ask whether a reasonable consumer could read the deployment as replacing an identifiable human role. If the answer is yes, the function enters a higher risk tier regardless of its technical category or execution quality.

Calibrate to audience composition before deployment. Aggregate consumer sentiment scores mask how strongly audience composition shapes bottom-tier outcomes. For creative advertising and virtual spokesperson functions, identify the audience segments most likely to react strongly to AI use, whether those segments are defined by product engagement, prior AI attitudes, or other category-specific factors, and plan for a concentrated negative reaction even when average sentiment looks acceptable.

Treat disclosure as a design constraint. The EU AI Act's consumer-facing disclosure obligations take effect in August 2026, and disclosure volume across all four functions will rise regardless of brand preference. Because identical disclosure produced sharply different consumer responses across functions, a function's position in the two-tier structure should drive both how much to invest in it and how to execute the disclosure. Bottom-tier deployments require the additional design work to preserve the human judgment signals that would otherwise be displaced.

Function-Specific Checklists

AI Creative Advertising

Deployment readiness

- AI is effective for concepting, layout exploration, and production acceleration. Keep AI in the production workflow rather than in the consumer-facing output.
- If business needs require AI-generated final assets, treat the deployment as bottom-tier risk and apply the remaining items in this checklist before launch.

Human element to preserve

- Retain human execution on the elements consumers use to infer judgment and taste: styling choices, visual composition, and product presentation.
- Ensure AI-generated product imagery can withstand disclosure without raising accuracy concerns.

Framing and disclosure

- When disclosing, specify what the AI did. For example, "AI-assisted texture enhancement" triggers less scrutiny than "AI-created advertisement."
- Avoid positioning AI as the creative author. The credibility cost concentrates on the elements where consumers expect human creative judgment.

AI Personalized Content Delivery

Deployment readiness

- Among the four functions tested under identical disclosure conditions, this one received the most favorable consumer response.

- Most personalized content delivery currently flows through platform-level recommendation engines that brands do not control. Identify where you can build direct brand-to-consumer personalization assets.

Framing and disclosure

- Frame the AI as working for the consumer. Language that emphasizes “based on your preferences” or “chosen for you” reinforces the utility reading that drives acceptance.
- Avoid framing that makes the data exchange salient. When consumers perceive personalization as surveillance rather than service, acceptance erodes. The minority who react negatively to this function react to the data collection, so proactive transparency about what data you use and why addresses the actual concern.

AI Customer Service Chatbots

Deployment readiness

- Chatbots in routine service contexts (password resets, order tracking, basic troubleshooting) operate in the top tier. Deploy with standard risk management.
- For high-stakes or emotionally complex interactions, maintain a clear human fallback path. Consumer acceptance of chatbots concentrates in routine scenarios.

Transition risk

- If chatbot deployment coincides with or is perceived as replacing human support staff, the function risks sliding from top tier to bottom tier.

Framing and disclosure

- Chatbot disclosure carries low cost in routine contexts. Plain disclosure is sufficient.
- Position the chatbot as supplementing human capacity. Consumer acceptance is grounded in passive tolerance of a familiar tool, and that tolerance depends on the chatbot being read as an addition.

AI Virtual Brand Spokespeople

Default position

- Under current consumer conditions, this function carries the highest deployment risk of the four. Four failure modes operate simultaneously: uncanny valley reactions, labor displacement concerns, brand authenticity challenges, and product representation doubts. No single design adjustment resolves all of them.

- Default to avoiding consumer-facing deployment of AI-generated brand personas unless a specific business case justifies the compounding risk.

If deployment is required

- Avoid replicating identifiable human performers. The Fortnite Darth Vader case demonstrates that replicating a recognizable voice generates both institutional backlash and alignment risk simultaneously.
- If using an LLM-driven persona in open consumer interaction, invest in behavioral guardrails before launch. Alignment failures in public-facing environments circulate immediately.

Reassessment

- Monitor consumer sentiment toward virtual personas at the industry level. If acceptance shifts, revisit deployment feasibility.

11. Limitations

This study has several constraints that affect how broadly the findings can be applied.

The consumer survey sample ($N = 150$) is smaller than the industry benchmarks it engages with, including the IAB and Sonata Insights study (505 consumers) and the USC Annenberg Global Communication Report (1,000+ professionals). The within-subjects design and the consistency of effects across all three outcome measures (all cross-tier gaps significant at $p < .001$) support internal validity. Generalizability to the broader U.S. consumer population is more limited.

Participants were recruited through Prolific, an online research panel. The sample skewed younger (67% between ages 25 and 44) and more educated (60% holding a bachelor's degree or higher) than the U.S. adult population. The sample's baseline familiarity with AI was also high: 49% described themselves as very or extremely familiar with AI in marketing, and 95% had previously used a customer service chatbot. The two-tier pattern identified here emerged among relatively AI-literate consumers. Whether the same pattern holds among less familiar populations remains an open question.

The survey used scenario-based descriptions rather than live interactions with actual AI systems. Participants read about each AI function and responded to that description. Real-world encounters with a chatbot, a personalized recommendation engine, or a virtual spokesperson may produce different reactions than reading about them. The scenarios also used a single fictional clothing brand across all four functions. Consumer

responses to AI may vary by product category, and findings from a clothing context may not transfer directly to categories like financial services, healthcare, or food.

The survey captures a single point in time. Consumer attitudes toward AI in marketing are shifting rapidly. SurveyMonkey (2026) found that the share of U.S. adults expressing concern about AI in everyday life rose from 27% to 33% in one year. The two-tier structure documented here reflects consumer sentiment during the survey window and may not remain stable as exposure increases and regulatory conditions change.

As noted in Section 9.1, the survey did not include a behavioral measure of prior experience with personalized recommendation systems. The familiarity analysis for that function relies on the general AI marketing familiarity score rather than a direct indicator. This limits the strength of the familiarity ruling for personalized content specifically.

The practitioner evidence draws on a single formal interview, field observation at one organization, and solicited online responses. The interview and field observation are both situated in the gaming sector, and while the solicited responses introduced cases from other industries, the evidence base remains narrow. The patterns identified in Section 8, including the concepting norm, the platform-mediated structure of personalized content, and the audience split between core and mass-market consumers, may reflect industry-specific dynamics that do not generalize to other sectors. The practitioner findings should be read as illustrative cases that generate hypotheses for further research rather than as representative evidence of industry-wide practice.

12. Conclusion

The problem this paper opened with is a measurement gap. Practitioners are deploying AI across consumer-facing marketing functions at accelerating rates. Consumer sentiment research exists for individual functions in isolation. No primary study had compared consumer responses to all four functions side by side, under controlled conditions, using the same participants and the same disclosure language. Without that comparison, practitioners had no way to know whether the resistance they observed in one function was specific to that function or a general consumer reaction to AI in marketing.

The consumer survey answered that question. The resistance is function-specific. Identical AI disclosure produced sharply different consumer responses depending on which function the AI occupied. Personalized content delivery and customer service chatbots scored at or above the scale midpoint across brand attitude, perceived authenticity, and purchase intention. Creative advertising and virtual spokespeople scored below it. The gap between tiers was significant at $p < .001$ on every measure. The gap within each tier was not.

The theoretical explanation centers on source type. When consumers perceive AI as a functional system operating in the background, they evaluate it on whether it works. When they perceive it as occupying a role previously held by a human, they evaluate it on whether it deserves that role. The same disclosure triggers different evaluative standards depending on the function it describes. This is the mechanism that produces the two-tier split, and it held across high-familiarity and low-familiarity participants alike.

The practitioner evidence confirmed the consumer pattern through independent observation. The consistent trigger for backlash across all practitioner sources was replacement framing, the same boundary the survey data identified through source attribution theory. The two data streams also revealed an industry workflow adopted for legal reasons that accidentally preserved the human judgment signals consumers use to evaluate credibility, and a structural gap in which the function consumers accept most readily is the one brands control least.

For practitioners reading this paper, the core takeaway is that AI deployment decisions should be made function by function. The four functions examined here do not carry equivalent risk, and treating them as a single category produces decisions calibrated to the wrong level of resistance. The recommendations in Chapter 10 translate this finding into deployment judgments organized by function and by the cross-cutting principles that emerged from both data streams.

The study has clear limitations in sample size, sector concentration, and the cross-sectional nature of the data. Consumer attitudes toward AI are moving. What this paper offers is a controlled baseline for a comparison that did not previously exist, and a structure for updating that comparison as conditions change.

Bibliography

Beverland, M. B., & Farrelly, F. J. (2010). The quest for authenticity in consumption:

Consumers' purposive choice of authentic cues to shape experienced outcomes.

Journal of Consumer Research, 36(5), 838–856

Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative*

research in psychology, 3(2), 77-101

Bui, H. T. (2025). Examining the effect of AI advertising involvement disclosure on

advertising value and purchase intentions. *Journal of Research in Interactive*

Marketing, 1–20. <https://doi.org/10.1108/JRIM-02-2025-0066>

Chen, Y., Wang, H., Rao Hill, S., & Li, B. (2024). Consumer attitudes toward AI-

generated ads: Appeal types, self-efficacy and AI's social role. *Journal of*

Business Research, 185, Article 114867.

<https://doi.org/10.1016/j.jbusres.2024.114867>

Cohen, J. (2013). *Statistical power analysis for the behavioral sciences*. Routledge.

Cronbach, L. J. (1951). Coefficient alpha and the internal structure of tests.

Psychometrika, 16(3), 297-334.

Crystal Dynamics. (2026). *Legacy of Kain: Defiance Remastered* [Video game]. Steam.

https://store.steampowered.com/app/3747730/Legacy_of_Kain_Defiance_Remastered/

- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm Aversion: People Erroneously Avoid Algorithms After Seeing Them Err. *Journal of Experimental Psychology. General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>
- Digiday. (2025, September 11). In Graphic Detail: Virtual influencers click with young audiences, yet brands' interest wanes. <https://digiday.com/media/in-graphic-detail-virtual-influencers-click-with-young-audiences-yet-brands-interest-wanes/>
- Dunn, O. J. (1961). Multiple comparisons among means. *Journal of the American statistical association*, 56(293), 52-64.
- EMARKETER. (2025, February 13). US consumer sentiment survey: AI adoption. <https://www.emarketer.com/content/us-consumer-sentiment-survey-ai-adoption>
- European Parliament and the Council of the European Union. (2024, June 13). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). EUR-Lex. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
- Franke, C., Groeppel-Klein, A., & Müller, K. (2023). Consumers' Responses to Virtual Influencers as Advertising Endorsers: Novel and Effective or Uncanny and Deceiving? *Journal of Advertising*, 52(4), 523–539. <https://doi.org/10.1080/00913367.2022.2154721>
- Friestad, M., & Wright, P. (1994). The Persuasion Knowledge Model: How People Cope with Persuasion Attempts. *The Journal of Consumer Research*, 21(1), 1–31. <https://doi.org/10.1086/209380>

Game Informer. (2025, May 19). SAG-AFTRA files unfair labor practice complaint against Epic Games due to A.I. Darth Vader. Game Informer.

<https://gameinformer.com/2025/05/19/sag-aftra-files-unfair-labor-practice-complaint-against-epic-games-due-to-ai-darth-vader>

Gartner. (2022, July 27). Gartner predicts chatbots will become a primary customer service channel within five years. <https://www.gartner.com/en/newsroom/press-releases/2022-07-27-gartner-predicts-chatbots-will-become-a-primary-customer-service-channel-within-five-years>

Gartner. (2024, July 9). Gartner survey finds 64 percent of customers would prefer that companies didn't use AI for customer service.

<https://www.gartner.com/en/newsroom/press-releases/2024-07-09-gartner-survey-finds-64-percent-of-customers-would-prefer-that-companies-didnt-use-ai-for-customer-service>

Gartner. (2026, January 15). Gartner predicts 60% of brands will use agentic AI to deliver streamlined one-to-one interactions by 2028 [Press release].

<https://www.gartner.com/en/newsroom/press-releases/2026-01-15-gartner-predicts-60-percent-of-brands-will-use-agentic-ai-to-deliver-streamlined-one-to-one-interactions-by-2028>

Greenhouse, S. W., & Geisser, S. (1959). On methods in the analysis of profile data. *Psychometrika*, 24(2), 95-112.

- Grigsby, J. L., Michelsen, M., & Zamudio, C. (2025). Service ads in the era of generative AI: Disclosures, trust, and intangibility. *Journal of Retailing and Consumer Services*, 84, 104231. <https://doi.org/10.1016/j.jretconser.2025.104231>
- Guest, G., Bunce, A., & Johnson, L. (2006). How Many Interviews Are Enough?: An Experiment with Data Saturation and Variability. *Field Methods*, 18(1), 59–82. <https://doi.org/10.1177/1525822X05279903>
- Hardcastle, K., Vorster, L., & Brown, D. M. (2025). Understanding Customer Responses to AI-Driven Personalized Journeys: Impacts on the Customer Experience. *Journal of Advertising*, 54(2), 176–195. <https://doi.org/10.1080/00913367.2025.2460985>
- Horvath, B. (2024, November 18). Coca-Cola causes controversy with AI-made ad. NBC News. <https://www.nbcnews.com/tech/innovation/coca-cola-causes-controversy-ai-made-ad-rcna180665>
- Hovland, C. I., Janis, I. L., & Kelley, H. H. (1953). *Communication and persuasion*.
- Ingram, M. B. (2023, October 31). The Finals responds to AI criticism. *Game Rant*. <https://gamerant.com/the-finals-beta-criticized-ai-voices-voice-acting/>
- Interactive Advertising Bureau. (2026, January 15). AI transparency and disclosure framework. <https://www.iab.com/guidelines/ai-transparency-and-disclosure-framework>

Interactive Advertising Bureau & Sonata Insights. (2024, December). The AI ad gap: Why young consumers aren't yet buying into gen AI ads. IAB.

<https://www.iab.com/insights/why-young-consumers-avoid-gen-ai-ads/>

Interactive Advertising Bureau & Sonata Insights. (2026, January). The AI ad gap widens. IAB. <https://www.iab.com/insights/the-ai-gap-widens/>

Kim, D., & Wang, Z. (2024). Social media influencer vs. virtual influencer: The mediating role of source credibility and authenticity in advertising effectiveness within AI influencer marketing. *Computers in Human Behavior: Artificial Humans*, 2(2), Article 100100. <https://doi.org/10.1016/j.chbah.2024.100100>

Koles, B., Audrezet, A., Moulard, J. G., Ameen, N., & McKenna, B. (2024). The authentic virtual influencer: Authenticity manifestations in the metaverse. *Journal of Business Research*, 170, Article 114325. <https://doi.org/10.1016/j.jbusres.2023.114325>

Koning, B., & Voorveld, H. A. M. (2025). Disclaimer! This Content Is AI-Generated: How AI-Disclosures Influence Trust in Advertisements and Organizations. *Journal of Interactive Advertising*, 25(3), 240–253. <https://doi.org/10.1080/15252019.2025.2554149>

Lee, H., Shin, M., Yang, J., & Chock, T. M. (2025). Virtual Influencers vs. Human Influencers in the Context of Influencer Marketing: The Moderating Role of Machine Heuristic on Perceived Authenticity of Influencers. *International Journal of Human-Computer Interaction*, 41(10), 6029–6046. <https://doi.org/10.1080/10447318.2024.2374100>

- Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational Behavior and Human Decision Processes*, 151, 90–103. <https://doi.org/10.1016/j.obhdp.2018.12.005>
- Luo, X., Tong, S., Fang, Z., & Qu, Z. (2019). Frontiers: Machines vs. Humans: The Impact of Artificial Intelligence Chatbot Disclosure on Customer Purchases. *Marketing Science (Providence, R.I.)*, 38(6), Article mksc. 2019.1192. <https://doi.org/10.1287/mksc.2019.1192>
- Mauchly, J. W. (1940). Significance test for sphericity of a normal n-variate distribution. *The annals of mathematical statistics*, 11(2), 204-209.
- McAloon, A. (2019). Steam embraces machine learning with interactive recommendation tool. *Game Developer*. <https://www.gamedeveloper.com/business/steam-embraces-machine-learning-with-interactive-recommendation-tool>
- McKinsey & Company. (2025, November). The state of AI in 2025: Agents, innovation, and transformation. https://www.mckinsey.com/~media/mckinsey/business%20functions/quantumblack/our%20insights/the%20state%20of%20ai/november%202025/the-state-of-ai-2025-agents-innovation_cmyk-v1.pdf
- Meltwater. (2025, June). The Rise of AI Influencers: What the Data Says About Consumer Sentiment and Interest. <https://www.meltwater.com/en/blog/ai-influencers-consumer-data>

Middler, J. (2023, October 31). The Finals studio responds to backlash at its decision to use AI voice acting. Video Games Chronicle.

<https://www.videogameschronicle.com/news/the-finals-studio-responds-to-backlash-at-its-decision-to-use-ai-voice-acting/>

Mordor Intelligence. (2026). Chatbot market size & share analysis: Growth trends and forecast (2026–2031). <https://www.mordorintelligence.com/industry-reports/global-chatbot-market>

Muck Rack. (2025). The state of PR 2025. Muck Rack.

<https://muckrack.com/research/state-of-pr>

NielsenIQ. (2024, December 12). NIQ research uncovers hidden consumer attitudes toward AI-generated ads. NielsenIQ. <https://nielseniq.com/global/en/news-center/2024/niq-research-uncovers-hidden-consumer-attitudes-toward-ai-generated-ads>

PinkNews. (2025, May 21). Fortnite issues fix after players make AI-powered Darth Vader use homophobic slurs. PinkNews.

<https://www.thepinknews.com/2025/05/21/fortnite-darth-vader-ai-slurs/>

Qin, X., Zhou, X., Chen, C., Wu, D., Zhou, H., Dong, X., Cao, L., & Lu, J. G. (2025). AI Aversion or Appreciation? A Capability-Personalization Framework and a Meta-Analytic Review. *Psychological Bulletin*, 151(5), 580–599.

<https://doi.org/10.1037/bul0000477>

Qiu, X., Wang, Y., Zeng, Y., & Cong, R. (2025). Artificial Intelligence Disclosure in Cause-Related Marketing: A Persuasion Knowledge Perspective. *Journal of*

Theoretical and Applied Electronic Commerce Research, 20(3), Article 193.

<https://doi.org/10.3390/jtaer20030193>

Randall, H. (2023, October 30). The Finals uses AI text-to-speech because it can produce lines "in just a matter of hours rather than months," baffles actual voice actors. PC Gamer. <https://www.pcgamer.com/the-finals-uses-ai-text-to-speech-because-it-can-produce-lines-in-just-a-matter-of-hours-rather-than-months-baffles-actual-voice-actors/>

Reich, T., Kaju, A., & Maglio, S. J. (2023). How to overcome algorithm aversion: Learning from mistakes. *Journal of Consumer Psychology*, 33(2), 285–302. <https://doi.org/10.1002/jcpy.1313>

Salesforce. (2024). State of the AI connected customer: 7th edition.

<https://www.salesforce.com/en-us/wp-content/uploads/sites/4/documents/research/State-of-the-Connected-Customer.pdf>

Shapiro, S. S., & Wilk, M. B. (1965). An analysis of variance test for normality (complete samples). *Biometrika*, 52(3-4), 591-611.

Shavelson, R., & Steedle, J. (2022). Coefficient Alpha and the Internal Structure of Tests. In B. B. Frey (Ed.), *The SAGE Encyclopedia of Research Design* (Vol. 4, pp. 213–215). SAGE Publications, Inc. <https://doi.org/10.4135/9781071812082.n86>

Shoenberger, H., & Kim, E. (Anna). (2023). Explaining purchase intent via expressed reasons to follow an influencer, perceived homophily, and perceived authenticity.

International Journal of Advertising, 42(2), 368–383.

<https://doi.org/10.1080/02650487.2022.2075636>

Spears, N., & Singh, S. N. (2004). Measuring Attitude toward the Brand and Purchase Intentions. *Journal of Current Issues and Research in Advertising*, 26(2), 53–66.

<https://doi.org/10.1080/10641734.2004.10505164>

Sprout Social. (2025). The rise of virtual influencers: are they here to stay?.

<https://sproutsocial.com/insights/virtual-influencers/>

Statista. (2022). Share of consumers who bought or would be unlikely to buy products promoted by virtual influencers in the United States as of 2022.

<https://www.statista.com/statistics/1300319/consumers-bought-products-promoted-virtual-influencers-us/>

Straits Research. (2025). Virtual influencer market size, share & trends analysis report.

<https://straitsresearch.com/report/virtual-influencer-market>

Sundar, S. S., & Nass, C. (2001). Conceptualizing sources in online news. *Journal of Communication*, 51(1), 52–72. [https://doi.org/10.1111/j.1460-](https://doi.org/10.1111/j.1460-2466.2001.tb02872.x)

[2466.2001.tb02872.x](https://doi.org/10.1111/j.1460-2466.2001.tb02872.x)

SurveyMonkey. (2026). AI sentiment study.

<https://www.surveymonkey.com/curiosity/surveymonkey-research-ai-sentiment-study/>

Trope, Y., & Liberman, N. (2010). Construal-Level Theory of Psychological Distance.

Psychological Review, 117(2), 440–463. <https://doi.org/10.1037/a0018963>

USC Center for Public Relations. (2025). Mind the gap: 2025 global communication report. USC Annenberg School for Communication and Journalism.
<https://annenberg.usc.edu/research/center-public-relations/global-communication-report>

Watson, J., Valsesia, F., & Segal, S. (2024). Assessing AI receptivity through a persuasion knowledge lens. *Current opinion in psychology*, 58, 101834.
<https://doi.org/10.1016/j.copsyc.2024.101834>

Wortel, C., Vanwesenbeeck, I., & Tomas, F. (2024). Made with Artificial Intelligence: The Effect of Artificial Intelligence Disclosures in Instagram Advertisements on Consumer Attitudes. *Emerging Media (Online)*, 2(3), 547–570.
<https://doi.org/10.1177/27523543241292096>

Zendesk. (2024). CX trends 2026: Unlock the power of intelligent CX.
<https://www.zendesk.com/blog/cx-trends-2024>

Primary Data Sources:

Consumer Survey

Zheng, H. (2026). AI in brand marketing: A study on consumer perceptions [Unpublished survey data]. USC Annenberg School for Communication and Journalism. Administered via Qualtrics; participants recruited through Prolific (N = 150).

Practitioner Interview

Fowler, S. (2026, March). Interview on AI deployment in gaming marketing [Recorded interview]. USC Annenberg DSM 596 Capstone Research. Conducted by H. Zheng.

Field Observation

Zheng, H. (2026, April 7). Field observation notes: AI marketing practices at a mobile gaming publisher [Unpublished field notes]. USC Annenberg DSM 596 Capstone Research.

Reddit Solicited Commentary

Zheng, H. (2026, April). Has AI in marketing ever triggered real consumer backlash? How did brands handle it? [Post]. r/AskMarketing, Reddit.
<https://www.reddit.com/r/AskMarketing/comments/1s9q5y7/>

Zheng, H. (2026, April). Brands using AI in marketing without telling anyone: Is this becoming the norm? [Post]. r/AskMarketing, Reddit.
<https://www.reddit.com/r/AskMarketing/comments/1sehrp1/>