

REVIEW ARTICLE OPEN ACCESS

Agent-Related Perceptions of Automated Versus Human Agents and the Role of Task-Related Risk: A Meta-Analysis

Sandra Miederer¹  | Katja Gelbrich¹  | Holger Roschk^{2,3}  | Masoumeh Hosseinpour⁴ 

¹Catholic University Eichstaett-Ingolstadt, Ingolstadt, Germany | ²Current: Aalborg University Business School, Aalborg University, Aalborg, Denmark | ³Incoming: Department of Marketing, BI Norwegian Business School, Oslo, Norway | ⁴Aarhus University, Aarhus C, Denmark

Correspondence: Sandra Miederer (SMiederer@ku.de)

Received: 2 October 2025 | **Revised:** 19 June 2026 | **Accepted:** 3 July 2026

Keywords: agent-related perceptions | automated agents | customer responses | digital agents | path model | robots | task-related risk

ABSTRACT

Firms often seek to replace human agents (HAs) with automated agents (AAs) for customer service, but customers still tend to react more negatively to AAs than to HAs. This meta-analysis examines agent-related perceptions (i.e., humanlikeness, warmth, and competence), which represent the key upstream customer reactions shaping downstream responses. Specifically, it clarifies the hierarchy of these perceptions and examines the moderating effect of task-related risks (i.e., functional, psychological, and social) on these perceptions. Meta-analytic path modeling with subgroup analyses, including 1191 effect sizes, provides novel insights. First, there is indication that perceived humanlikeness is not a prerequisite for perceived warmth and competence; all three perceptions emerge in parallel. Second, tasks entailing functional or psychological risk largely weaken the negative perceptions of AAs (vs. HAs). For psychologically risky tasks, AAs are not perceived as less competent than HAs. Tasks entailing social risk largely strengthen the negative perceptions of AAs (vs. HAs). These findings challenge two prevailing assumptions in the services marketing literature: that customers favor humans over AAs simply because they are not humanlike and that risk generally discourages AA usage. The results have important managerial and future research implications.

1 | Introduction

Advancements in technology have driven the rise of automated agents (AAs). These are system-based autonomous entities, either physical (e.g., Pepper) or digital (e.g., ChatGPT) (Blut et al. 2021), that interact with customers to perform various services (Wirtz et al. 2018), such as selling products (Holthöwer and van Doorn 2023). Using AAs can increase the speed and efficiency of customer service (Wirtz et al. 2018). Simultaneously, AAs may simulate a social presence, prompting perceptions of warmth and competence as immediate customer reactions shaping downstream responses (van Doorn et al. 2017). As these perceptions originally apply to human interaction partners (Fiske et al. 2007), it is further argued that perceived warmth and competence are reinforced by AAs'

perceived humanlikeness, that is, the anthropomorphization of AAs (van Doorn et al. 2017).

This notion of AAs as “efficient social actors” has fueled multiple studies examining whether they can replace human agents (HAs) by comparing customer responses to both agent types (e.g., Sheng et al. 2025). Two meta-analyses have consolidated this knowledge (Gelbrich et al. 2026; Zehnle et al. 2025). Both support the argument that customers are still less familiar with AA-based interactions (van Doorn et al. 2025) and thus react more negatively to them than to HAs. Yet, AAs' inferiority is less pronounced than assumed and mainly occurs at the emotional and perceptual level. Furthermore, the negative effect of AAs (vs. HAs) on customer responses is subject to contingencies (e.g., task, context) that reduce AAs' relative inferiority.

This is an open access article under the terms of the [Creative Commons Attribution](https://creativecommons.org/licenses/by/4.0/) License, which permits use, distribution and reproduction in any medium, provided the original work is properly cited.

© 2026 The Author(s). *Psychology & Marketing* published by Wiley Periodicals LLC.

Yet, the meta-analyses reveal two gaps. First, both examine multiple customer responses but do not test their order, for which opposing rationales exist. On the one hand, the notion of AAs as having a social presence, akin to humans, implies a hierarchical structure: increasing AAs' perceived humanlikeness fuels warmth and competence (van Doorn et al. 2025). However, it can be argued that AAs have superior machine characteristics such as speed and reliability (Wirtz et al. 2018) and thus do not have to be perceived as humanlike to be suitable substitutes for HAs (Gelbrich et al. 2026). The latter perspective implies that perceived warmth and competence are independent of humanlikeness, speaking to a parallel structure. Single studies show inconsistent results, either supporting a hierarchical (Im et al. 2023) or a parallel structure (Belanche et al. 2021). *Hence, our first research question is: Do agent-related perceptions of AAs (vs. HAs) occur in parallel, or does humanlikeness precede warmth and competence?*

Second, none of the two meta-analyses includes risk as a contingency. This is surprising, because perceived risk is a psychological construct that shapes purchase decisions (Baker et al. 2016). This includes the use of new technologies such as AAs (Belanche et al. 2025). Consolidated research on AA adoption alone—that is, without AA-HA comparison—shows that the AA-inherent perceived risk inhibits their usage (e.g., Mehta et al. 2022). However, AAs replace humans in performing tasks for customers (Wirtz et al. 2018), and tasks inherently carry specific risks (Murray and Schlacter 1990). For example, private tutoring involves psychological risk, as students' self-esteem may suffer after incorrect responses. Thus, when comparing AAs versus HAs, the focus shifts from AA-inherent to task-related risk. Both agent types have unique characteristics that may be perceived as more or less appropriate for handling specific task-related risks. *Hence, our second research question is: Do task-related risks influence the negative effect of AAs (vs. HAs) on agent-related perceptions?*

To answer these research questions, we present a conceptual framework of customer responses to AAs versus HAs in customer service, tested in a meta-analysis with 1191 effect sizes. For our framework, we take the sequence “AAs (vs. HAs) → perceptions → appraisals → intentions” as a starting point and refine it in two ways. First, we test competing path models for the structure of agent-related perceptions: a hierarchical model (i.e., perceived humanlikeness precedes warmth and competence) and a parallel model (i.e., agent-related perceptions are directly caused by AAs vs. HAs). Second, we test three risk types (i.e., functional, psychological, and social) as moderators for the effect of AAs (vs. HAs) on agent-related perceptions.

Our study contributes to the services marketing literature in two ways. First, we clarify the order of agent-related perceptions. This matters because the two competing models provide different explanations for AAs' inferiority to HAs and how to change it. A hierarchical order implies that AAs are inferior because they are not human(like), and thus less warm and competent. Our results indicate a parallel order, suggesting that AAs may be perceived as warm and competent even when not perceived as humanlike. This reshapes the current understanding of social perception in AA-based customer interactions as a process that is inherently linked to anthropomorphization

(van Doorn et al. 2025), fueling customers' perceptions of AAs as warm and intelligent (Blut et al. 2021). We provide initial evidence that anthropomorphization may occur but is possibly not necessary.

Second, we shed light on the occurrence of these agent-related perceptions by showing that perceived risk plays a nuanced role in this process. Prior AA adoption research finds that risk generally hinders AA usage, without specifying the role of agent-related perceptions (Mehta et al. 2022). Social psychological research suggests that interactions with nonhuman social actors entail high uncertainty, which may be reduced by assigning human qualities to them (Epley et al. 2007; Epley and Waytz 2010). This allows for the possibility of positive AA perceptions even under high risk. We propose that agent-related perceptions depend on the risk type entailed in a particular task. Considering task-related risk as a multifaceted construct (Derbaix 1983), the results suggest that task-related functional and psychological risks may indeed work for AAs: Both largely weaken the negative effect of AAs (vs. HAs) on perceived humanlikeness, warmth, and competence. Only social risk may work against AAs, largely strengthening their inferior perceptions.

Our results provide guidance for practitioners. Since perceived humanlikeness, warmth, and competence occur in parallel, service managers may not need costly humanlike design cues. Instead, they can make AAs warmer and more competent. Furthermore, managers should use AAs for functionally or psychologically risky tasks, but rely on HAs for socially risky tasks.

2 | Conceptual Development

2.1 | Overview of the Conceptual Framework

Figure 1 displays our meta-analytic framework. The customer responses include, in sequence, agent-related perceptions, agent- and firm-related customer appraisals, and firm-related customer intentions (Ajzen 1991; Belanche et al. 2021; Gelbrich et al. 2026; Wentzel 2009). *Agent-related perceptions* refer to perceived humanlikeness as the degree to which an agent is ascribed humanlike emotions, characteristics, or motivations, perceived warmth as the degree to which an agent is perceived as having traits related to positive intentions, and perceived competence as the degree to which an agent is perceived as having traits related to capability (Gelbrich et al. 2026). *Customer appraisals* pertain to assessments of the agent or the firm based on associations linked to them (e.g., trust in the agent, satisfaction with the firm). *Firm-related customer intentions* refer to the propensity to perform positive behaviors toward the firm (Hoyer et al. 2024). Web Appendix A provides typical aliases.

Using this framework, we test whether agent-related perceptions occur in a parallel or hierarchical order and whether these perceptions are risk-dependent. Regarding the order of perceptions, we present the concept of social perception (Fiske et al. 2007), which suggests a hierarchical order, and the theory of mind perception, in which we ground our hypothesis and which suggests a parallel order (H1) (Gray et al. 2007). We draw on the concept of task-related risk (Murray and Schlacter 1990;

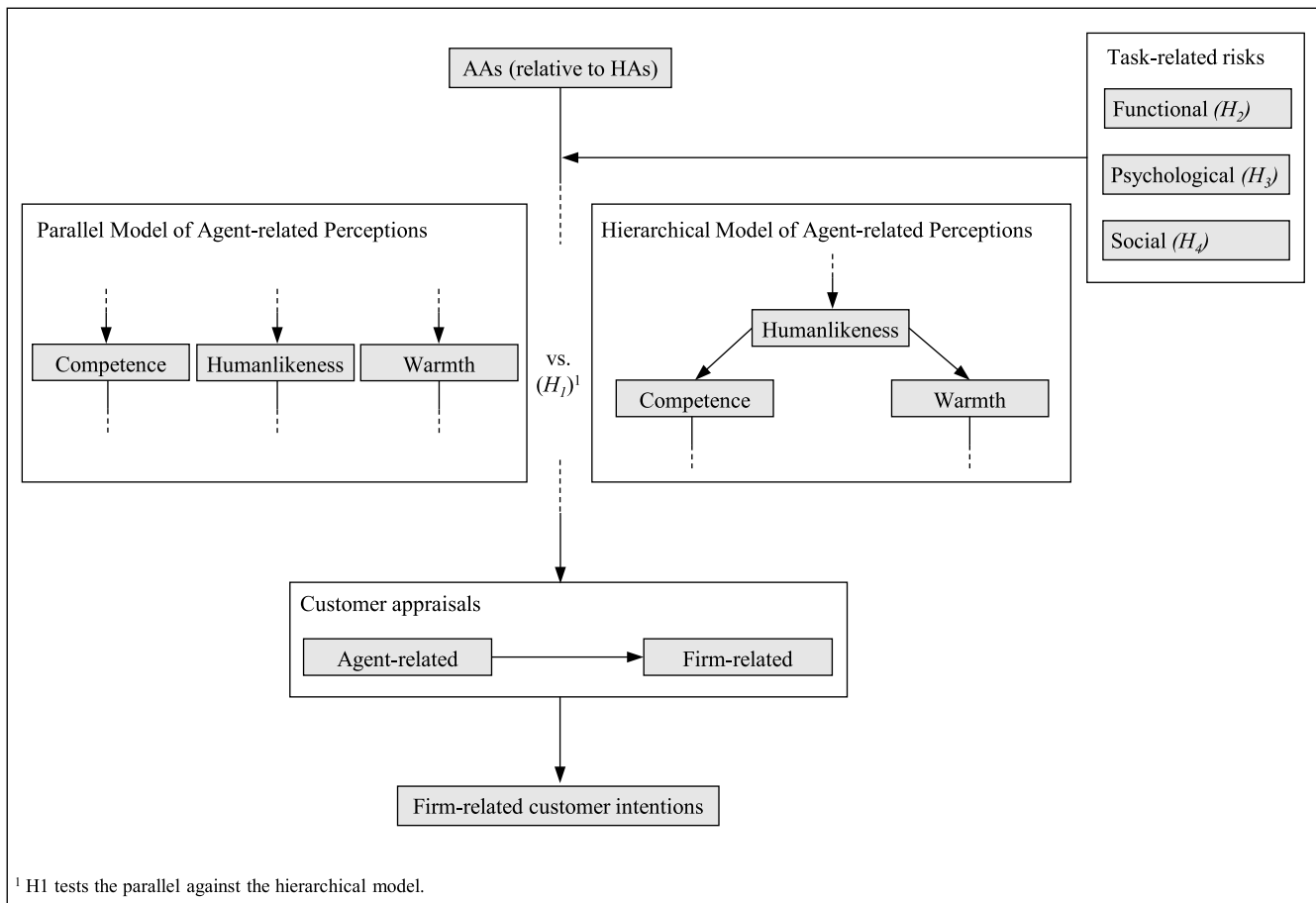


FIGURE 1 | Meta-analytic framework.

Sohn 2024) to hypothesize moderating effects of risk types on agent-related perceptions (H2–H4).

2.2 | Agent-Related Perceptions

2.2.1 | Hierarchical Model

A hierarchy of agent-related perceptions can be inferred from the concept of social perception, suggesting that people assess human interaction partners on two dimensions. Perceived warmth refers to traits associated with positive intentions, such as friendliness. Perceived competence refers to traits reflecting the ability to act on those intentions, such as skillfulness (Fiske et al. 2007). This assessment of interaction partners has been transferred to AAs (Heerink et al. 2008), as AAs can prompt customers to perceive an (automated) social presence (i.e., company of another person). Accordingly, customers tend to consider AAs as social actors, triggering social perceptions of warmth and competence (van Doorn et al. 2017). Furthermore, these perceptions may be reinforced by AAs' perceived humanlikeness, also known as anthropomorphization (van Doorn et al. 2017). It describes the tendency to imbue nonhuman objects with humanlike emotions, characteristics, or motivations (Epley et al. 2007). Perceived humanlikeness can be triggered by human cues in, for example, AAs' appearance (Belanche et al. 2021) or voice (Im et al. 2023), and thus may reinforce perceptions of warmth and competence, usually assigned to HAs. This suggests a hierarchical order where perceived humanlikeness precedes warmth and competence

(Im et al. 2023). However, we consider a parallel structure as more reasonable.

2.2.2 | Parallel Model

A parallel order of agent-related perceptions can be inferred from the theory of mind perception (Gray et al. 2007). This theory suggests that people seek to read other agents' minds along two dimensions. Experience refers to the perceived capacity for sensation and feeling and is related to warmth; agency refers to the perceived capacity for reasoned action and is related to competence (Epley and Waytz 2010; Gray et al. 2007). These inferences help people understand and predict social agents' communication and coordinate their own behavior accordingly. As such, the theory of mind perception is similar to, and includes social perception (Epley and Waytz 2010). Likewise, it has been applied to AAs, arguing that they fall short of HAs in this regard (Srinivasan and Sarial-Abi 2021). However, this theory differs from the concept of social perception in two important ways.

First, from the outset, it widens the notion of social agents to any entity that interacts with others, including animals, supernatural authorities such as God, or technology (Gray et al. 2007), which is an anthropomorphization (Epley and Waytz 2010). Second, it presents the conditions that activate mind reading: interdependence between us and the other agents, powerful agents, or surveillance agents that evaluate us,

making it important to anticipate and influence their impressions (Epley and Waytz 2010). As such, this theory acknowledges anthropomorphization (i.e., general humanlikeness perceptions) as a natural response to self-propelled nonhuman agents such as AAs. Importantly, it presents mind perception (i.e., warmth and competence) as a distinct psychological mechanism that is also triggered by such agents as soon as they meet the aforementioned conditions, without requiring a general person perception (Epley and Waytz 2010). All three conditions apply to AAs: they interact with customers (e.g., selling products in pharmacies), have power (e.g., guiding the way at airports), or evaluate customers (e.g., deciding upon loan approvals). Hence, customers need to infer AAs' positive intentions and abilities, for example, whether AAs recommend the best or most profitable product (i.e., positive vs. negative intention; warmth), and whether they have or have not been trained with sufficient data ensuring good advice (i.e., high vs. low ability; competence).

In summary, these arguments suggest that both humanlikeness and social perceptions are caused by the same stimulus (i.e., an AA replacing an HA), speaking to a parallel structure of perceived humanlikeness, warmth, and competence (Belanche et al. 2021). Thus:

H1. *Agent-related perceptions (i.e., humanlikeness, warmth, and competence) occur in a parallel rather than a hierarchical order.*

2.3 | Task-Related Risk Shaping Agent-Related Perceptions

Risk refers to the potential loss associated with a purchase decision and can be classified into three types: functional, psychological, and social (Sohn 2024; Taylor and Taylor 1974). While research has identified more categories, they originate in an incorrect outcome and can be subsumed under functional risk (Derbaix 1983). Since AAs' purpose is to perform tasks for customers (Wirtz et al. 2018), we refer to task-related risk, defining it as whether or not the task performed by an agent entails the likelihood of negative consequences. We hypothesize that risk types moderate the negative effect of AAs (vs. HAs) on the three agent-related perception because AAs may manage task-related risks better or worse than HAs, weakening or strengthening their perceived inferiority.

Task-related functional risk refers to whether the agents' task performance entails potential underperformance compared to customers' expectations (Sohn 2024), specifically outcome expectations (e.g., financial advice). These outcome expectations may be independent of the agent, as task accomplishment is agent-neutral per se. Thus, customers' (dis)confirmation of the agent's performance is primarily driven by the actual outcome level. This reasoning aligns with research suggesting that customers' preference for AAs depends on expected task outcomes (Clegg et al. 2024).

Whereas human performance inherently varies over time (Zeithaml et al. 2024), AAs are more capable of providing stable and correct outcomes due to their machine nature, which

qualifies them as reliable agents for task execution (Larkin et al. 2022). Thus, in settings with (vs. without) functional risk, AAs may better meet customer expectations, reducing their perceived inferiority relative to HAs. Still, differences across individual customers and tasks may make more variable outcomes acceptable (i.e., lower expectations). Thus, we consider our rationale to reflect a cross-study, or structural, tendency. Collectively, the post-decision affect from disconfirmation may translate into competence perceptions, where functional risk weakens the inferiority of AAs (vs. HAs). Furthermore, whereas AAs (vs. HAs) lack affectionate concern (Castelo et al. 2019), these attributes may be less relevant for outcome-focused disconfirmation, such that functional risk may reduce the inferiority of AAs also in relation to warmth and humanlikeness. Thus:

H2. *For tasks with functional risk, the negative effect of AAs (vs. HAs) on perceived (a) humanlikeness, (b) warmth, and (c) competence is weaker than for tasks without functional risk.*

Task-related psychological risk refers to whether agents' task performance (e.g., private teaching) entails potential harm to customers' psychological comfort and self-esteem (Hoyer et al. 2024; Sohn 2024). Two discomforting responses are feelings of judgment and embarrassment arising from the threat of unwanted evaluation by others (Pitardi et al. 2022). The prominent "other" involved is the AA or HA as customers' direct interaction partner; bystanders require a different approach, discussed below.

In contexts entailing psychological risks, customers may focus on maintaining their self-esteem (Pyszczynski et al. 2004), for which AAs' machine characteristics can be beneficial. AAs do not form opinions and cannot judge others (Holthöwer and van Doorn 2023). Hence, AAs do not threaten customers' self-esteem, for example, by disapproving of their inquiries, nor do they pose a threat of unwanted evaluation. Hence, in tasks with (vs. without) psychological risks, AAs' inferiority to HAs may be reduced in relation to customers' warmth (i.e., lacking hidden motives) and competence perceptions (i.e., handling customers' inquiries professionally) of AAs. Similarly, AAs' humanlike perceptions may benefit in such tasks, because being nonjudgmental is a desirable human trait (Nicol and de France 2024). Thus:

H3. *For tasks with psychological risk, the negative effect of AAs (vs. HAs) on perceived (a) humanlikeness, (b) warmth, and (c) competence is weaker than for tasks without psychological risk.*

Task-related social risk refers to inhibited task performance due to bystanders, resulting in status loss (Sohn 2024). Thus, social risk centers on the presence of bystanders. While agents act autonomously, their performance depends on how well people interact with them. Social facilitation theory proposes that bystanders decrease individuals' performance on complex or unfamiliar tasks (Bond and Titus 1983). Bystanders may further increase individuals' arousal, already heightened from a difficult task, resulting in overstimulation and poor performance (Bond and Titus 1983; Zajonc and Sales 1966).

AAs are employed in many settings with bystanders (e.g., check-in at a hotel lobby), and interacting with them remains unfamiliar and difficult for customers (van Doorn et al. 2025).

Furthermore, AAs (vs. HAs) have lower cognitive flexibility (Longoni et al. 2019), making it hard for them to effectively respond to all kinds of inquiries (Wirtz et al. 2018). Hence, bystanders may cause a suboptimal interaction of customers with AAs, exacerbated by AAs' cognitive inflexibility. Customers are likely to blame AAs, at least partly, for this suboptimal interaction. This tendency reflects a self-protection bias whereby people attribute negative events to external sources (Heider 1958), leading them to blame employees rather than themselves for failures (Bitner et al. 1994). Hence, customers are likely to perceive AAs less favorably, in terms of humanlikeness (i.e., a heightened sense of an alien interaction), warmth (i.e., lacking sensitivity to customers' stress), and competence (i.e., unable to properly respond to customers' input) that enhance AAs' inferiority. Thus:

H4. *For tasks with social risk, the negative effect of AAs (vs. HAs) on perceived (a) humanlikeness, (b) warmth, and (c) competence is stronger than for tasks without social risk.*

3 | Method

3.1 | Overview

We adapted Gelbrich et al.'s (2026) dataset for our purposes. Specifically, we updated it by adding studies published after June 2023. Furthermore, we included effect sizes among customer responses in our analyses, discarded responses with an insufficient number of correlations, and coded the three task-related risk types as moderators. For data analysis, we provided robust effect size estimates whenever possible. The resulting pooled effect sizes served to create correlation matrices for path modeling (H1) and subgroup analyses (H2–H4).

3.2 | Data Collection

3.2.1 | Literature Search

The search spans the period from 2000, when the topic first appeared, until November 2025. We searched for studies examining relationships between the core model constructs, namely (1) the effects of AAs (vs. HAs) on customer responses and (2) correlations among these responses (Figure 1). First, we screened three electronic databases (Web of Science, EBSCO, and ScienceDirect) with search terms related to AAs in general (e.g., AI) or to embodied and digital forms (e.g., robot, adaptive AI avatars), and combined them with terms indicating comparison (e.g., replacement/comparison and human/employee). Next, we screened relevant journals focusing on marketing (e.g., *Journal of Marketing*), service (e.g., *Journal of Service Research*), psychology (e.g., *Psychological Science*), human-machine interaction (e.g., *Computers in Human Behavior*), including interdisciplinary journals (e.g., *Psychology & Marketing*). The Internet was searched for gray literature (e.g., dissertations, conference proceedings), using platforms like SSRN and Google Scholar. Finally, backward and forward citation searches were conducted. When necessary, the authors of articles were contacted for additional data and unpublished work. We used clear inclusion criteria: AA versus HA comparison, employee-customer interaction, use of one of the above-mentioned customer responses, and sufficient data for effect size calculation.

3.2.2 | Sample

Database

The database consists of 162 articles, including 354 independent studies with 1193 retrieved effect sizes. Web Appendix B provides detailed information on how we altered and enriched Gelbrich et al.'s (2026) dataset. The average proportion of females is 51.0%; subjects' mean age is 33.5 years; they come from North America (44.6%), Asia (36.4%), Europe (15.3%), and Australia (3.7%). Web Appendices C and D list all included studies and articles.

Effect Size Retrieval

The effect sizes for the relationships between AAs (vs. HAs) and customer responses were culled from Gelbrich et al.'s (2026) database or newly added. All effect sizes were based on the standardized mean differences (Cohen's d), converted into Pearson's correlation coefficient r (Lipsey and Wilson 2001). Additionally, we retrieved the correlations between customer responses. When r was not reported in the study, we calculated it using the original dataset (if available) or converting β coefficients (Peterson and Brown 2005).

Outliers and Publication Bias

From the original database ($k = 1193$), we discarded two outliers, which were outside the interval of the mean plus or minus three times the standard deviation (Liadeli et al. 2023). When re-including them, the result of H1 remained stable. Results of H2–H4 were not affected because the outliers did not pertain to these effects. Furthermore, publication bias, assessed through fail-safe N , funnel plots, and effect size precision, was not a major problem (Web Appendix E).

3.2.3 | Coding Scheme

Table 1 lists the study variables. Specific attention was required for task-related risks. Each risk was coded as a separate variable (i.e., 0 = absence of risk type, 1 = presence of risk type). This allows us to capture multiple risks within a task, accounting for ambiguity from multiple risks. We used a binary scale to ensure clear categories, which is instrumental for subgroup analyses. The risk types (functional, psychological, and social) and some covariates (physical embodiment, real interaction, and field experiment) were subjective variables—coded by two independent coders, trained by the first author and unaware of the hypotheses. Given that all these variables were dichotomous, we used Cohen's Kappa for intercoder agreement. It ranged from 0.69 to 1.00, indicating substantial to perfect agreement (Landis and Koch 1977). Discrepancies were resolved through discussion until a final agreement was reached. The remaining objective variables were coded by the authors.

3.3 | Data Analysis

3.3.1 | Pooling Effect Sizes

We pooled the effect sizes r to achieve a meta-analytic correlation matrix for path modeling (Rubera and Kirca 2012). We followed common meta-analytical procedures: correcting r values for

TABLE 1 | Coding scheme.

Variable	Coding Scheme	Mean (SD)
Moderators		
Functional risk	The task performed by the agent [1] entailed potential underperformance compared to customers' expectations, [0] did not.	0.41 (0.52)
Psychological risk	The task performed by the agent [1] entailed potential harm to customers' psychological comfort and self-esteem, [0] did not.	0.15 (0.36)
Social risk	The task performed by the agent [1] entailed potential harm to customers' social status due to bystanders, [0] did not.	0.21 (0.54)
Covariates^a		
Proportion of females	Proportion of female participants in the sample	0.51 (0.18)
Mean age	Participants' mean age	33.49 (8.06)
Physical embodiment	The AA was a [1] physical agent, [0] digital agent.	0.30 (0.46)
Control variables	The analyses [1] included control variables, [0] did not.	0.06 (0.25)
Published work	The article [1] was published, [0] was not.	0.92 (0.27)
Top-tier journal ^b	The article was published in a [1] top-tier journal, [0] non-top-tier journal.	0.25 (0.43)
Business journal	The article was published in a [1] business journal, [0] non-business journal.	0.66 (0.47)
Real interaction ^c	The research design was based on [1/3] a real field experiment, [1/3] a reallab experiment, [−2/3] a scenario.	−0.37 (0.45)
Field experiment ^c	The research design was based on [1/2] a real field experiment, [−1/2] a real lab experiment, [0] a scenario.	0.04 (0.26)

^aWe used the same covariates as Gelbrich et al. (2026), with some exceptions. Physical embodiment was included because AAs can either be embodied (i.e., robots) or digital (i.e., virtual assistants). General population and crowd-based recruitment were removed due to multicollinearity with other covariates. For the effect of AAs (vs. HAs) on perceived humanlikeness, real interaction was also removed for the same reason.

^bThe journal impact factor (JIF) is used, with JIF ≥ 10 as the threshold for top-tier journals.

^cContrast-coded variables.

reliability, performing Fisher's z transformation, weighting effect sizes by their inverse variance, and running a random-effects multilevel model (see Gelbrich et al. 2026 for details).

For relationships between AAs (vs. HAs) and customer responses, we had sufficient effect sizes to adjust for covariates (Table 1). Specifically, we regressed the Fisher's z -transformed effect sizes on the covariates, again with an inverse variance weighting and a random-effects multilevel model (Web Appendix F). Then, we used the estimated regression coefficients to predict the effect size at mean covariate values, yielding the adjusted pooled effect size. This approach allowed robust estimation of the pooled effect size (Liadeli et al. 2023). For all analyses, we used the `rma.mv` function from the `metafor` package for RStudio (Viechtbauer 2010).

3.3.2 | Path Modeling (H1)

The pooled effect sizes across all model constructs represent the meta-analytical correlation matrix (Table 2). It was used for path modeling, along with the harmonic mean of the cumulative sample size for each relationship ($N=1315$; Rubera and Kirca 2012). For perceived humanlikeness, few correlations were available, including only one with firm-related customer intentions. We retained perceived humanlikeness because it is a core construct, and the single correlation was positive and moderate ($r=0.547$, $p<0.001$), supporting nomological and discriminant validity. Furthermore, a model without it showed an unchanged effect pattern. Path modeling was conducted using SPSS Amos (Version 28) with maximum-likelihood estimation (Rubera and Kirca 2012). To test H1, we compared the parallel model (i.e.,

humanlikeness, warmth, and competence are directly caused by AAs vs. HAs as a common stimulus) with the hierarchical model (i.e., perceived humanlikeness precedes warmth and competence, thus mediating the effect of AAs vs. HAs on these social perceptions).

3.3.3 | Subgroup Analyses (H2–H4)

Subgroup analyses were performed for the path model with the superior model fit (parallel model) to test H2–H4. For this purpose, we re-calculated the harmonic mean sample size and the effect sizes of AAs (vs. HAs) on perceived humanlikeness, warmth, and competence for both levels (0 and 1) of each risk moderator, using the same method as above. This resulted in two correlation matrices per moderator (Web Appendix G). Thus, risk types were treated independently; they only overlap for few effect sizes (Web Appendix H). The new correlation matrices were used in subgroup analyses for each moderator (i.e., functional, psychological, and social risk), using χ^2 -difference tests (Jöreskog 1971) with SPSS Amos (Version 28).

4 | Results

4.1 | Path Model Comparison

To test H1, we estimate a parallel and a hierarchical model. We let the error terms of the three perceptions (parallel model) or of warmth and competence (hierarchical model) correlate to account for the high correlations among these constructs. In the parallel model, we observe high correlations (r) between the

TABLE 2 | Correlation matrix.

	1	2	3	4	5	6	7
1. AAs (vs. HAs)	—	74 (12,831)	80 (15,482)	97 (19,817)	180 (37,221)	366 (83,719)	202 (54,961)
2. Perceived humanlikeness	−0.597*	—	5 (1031)	4 (641)	3 (306)	7 (1178)	1 (205)
3. Perceived warmth	−0.251*	0.752**	—	10 (3305)	13 (3641)	17 (3702)	13 (3371)
4. Perceived competence	−0.133*	0.566***	0.632***	—	9 (1547)	12 (2923)	5 (579)
5. Agent-related customer appraisals	−0.178*	0.658 [†]	0.666***	0.608***	—	23 (9558)	16 (4919)
6. Firm-related customer appraisals	−0.079*	0.628***	0.631***	0.666***	0.738***	—	54 (36,331)
7. Firm-related customer intentions	−0.134*	0.547***	0.575***	0.617**	0.698***	0.733***	—

Note: Sample-size-weighted mean correlations are displayed by off-diagonal entries in the lower-left triangle. Number of effect sizes and total sample sizes (in parentheses) are depicted by the off-diagonal entries in the upper-right triangle.

* $p \leq 0.05$; ** $p \leq 0.01$; *** $p \leq 0.001$.

[†] $p \leq 0.10$.

error terms: 0.775 (warmth and humanlikeness), 0.624 (warmth and competence), and 0.612 (humanlikeness and competence). In the hierarchical model, it is lower (0.380), as warmth and competence derive from humanlikeness. Letting error terms correlate is justified as there is theoretical or empirical indication of a higher-order construct (Fornell 1983). In the parallel model, the higher-order construct can be labeled as agent-related perceptions resulting from exposure to an AA, derived from theory of mind perception (Epley and Waytz 2010). In the hierarchical model, the higher-order construct refers to social perceptions (warmth and competence), which correlate according to prior research (Fiske et al. 2007). Furthermore, letting error terms correlate is justified, as the structural parameter estimates remain stable (Fornell 1983).

We first assume that the effect of AAs (vs. HAs) on firm-related customer intentions is fully mediated by perceptions and appraisals for both models, which yields unsatisfactory model fits (parallel model: $\chi^2[6] = 249.638$, $p < 0.001$, $\chi^2/df = 41.606$, SRMR = 0.047, RMSEA = 0.176; hierarchical model: $\chi^2[10] = 632.963$, $p < 0.001$, $\chi^2/df = 63.296$, SRMR = 0.089, RMSEA = 0.218). The modification indices indicate that this is due to direct paths (i.e., partial mediational structure), which we add in a second step. Figure 2 shows these models, with direct paths indicated by dotted lines. The parallel model yields a satisfactory fit ($\chi^2[3] = 7.894$, $p = 0.048$, $\chi^2/df = 2.631$, SRMR = 0.006, RMSEA = 0.035). The hierarchical model does not ($\chi^2[5] = 257.432$, $p < 0.001$, $\chi^2/df = 51.486$, SRMR = 0.063, RMSEA = 0.196), with χ^2/df and RMSEA exceeding the thresholds of 3.0 and 0.08, respectively (Imhof and Klaus 2020). These results favor the parallel model, supporting H1. We additionally examined two intermediate models (i.e., humanlikeness as an antecedent of warmth or competence); Web Appendix I reports their inferior fit. Path model comparisons across both levels of each risk moderator also support H1 (Web Appendix J). Web Appendix K provides evidence of discriminant validity for the favored parallel model.

The parallel model in Figure 2 reveals the following significant paths. AAs (vs. HAs) are perceived as less warm ($\beta = -0.251$, $p = 0.010$), less humanlike ($\beta = -0.597$, $p = 0.011$), and less competent ($\beta = -0.133$, $p = 0.011$). These perceptions are positively related to the agent-related customer appraisals (warmth: $\beta = 0.194$, $p = 0.020$; humanlikeness: $\beta = 0.504$, $p = 0.002$; competence: $\beta = 0.227$, $p = 0.006$) and firm-related customer appraisals (humanlikeness: $\beta = 0.409$, $p = 0.008$; competence: $\beta = 0.255$, $p = 0.008$). Both appraisals facilitate firm-related

customer intentions (agent-related customer appraisals: $\beta = 0.300$, $p = 0.008$; firm-related customer appraisals: $\beta = 0.401$, $p = 0.009$). There is a positive effect of agent-related customer appraisals on firm-related customer appraisals ($\beta = 0.375$, $p = 0.005$). As additional direct paths, AAs (vs. HAs) positively affect both appraisals (agent-related customer appraisals: $\beta = 0.202$, $p = 0.007$; firm-related customer appraisals: $\beta = 0.260$, $p = 0.023$) and perceived competence on firm-related customer intentions ($\beta = 0.167$, $p = 0.012$). The path from perceived warmth to firm-related customer appraisals is nonsignificant ($p = 0.289$).

All covariates are nonsignificant except for physical embodiment, which negatively affects warmth ($\beta = -0.217$, $p = 0.011$) and humanlikeness ($\beta = -0.425$, $p = 0.007$). This resonates with prior findings that physical AAs are less well received than digital AAs (Gelbrich et al. 2026), as humanlikeness is more salient for robots, making differences between AAs and HAs more apparent. Furthermore, metal bodies lack inherent warmth.

4.2 | Moderation Analyses

Table 3 presents the subgroup analyses results testing H2-H4. Regarding functional risk, the negative effect of AAs (vs. HAs) on perceived humanlikeness is weaker ($p < 0.001$) for tasks with functional risk ($\beta = -0.376$, $p = 0.009$) versus without it ($\beta = -0.603$, $p = 0.009$). Likewise, the negative effect on perceived competence is weaker ($p = 0.005$) for tasks with ($\beta = -0.057$, $p = 0.026$) versus without this risk ($\beta = -0.167$, $p = 0.011$). These results support H2a and H2c. The result for perceived warmth is nonsignificant ($p = 0.512$), so H2b is not supported.

Regarding psychological risk, the negative effect of AAs (vs. HAs) on perceived humanlikeness is weaker ($p < 0.001$) for tasks with ($\beta = -0.483$, $p = 0.010$) versus without this risk ($\beta = -0.641$, $p = 0.012$). Likewise, the negative effect on perceived warmth is weaker ($p = 0.047$) for tasks with ($\beta = -0.189$, $p = 0.015$) versus without this risk ($\beta = -0.266$, $p = 0.012$). The effect on perceived competence also differs ($p < 0.001$); it is nonsignificant for tasks with this risk ($\beta = 0.026$, $p = 0.431$) and negative for tasks without it ($\beta = -0.163$, $p = 0.014$). Thus, H3a-H3c are supported.

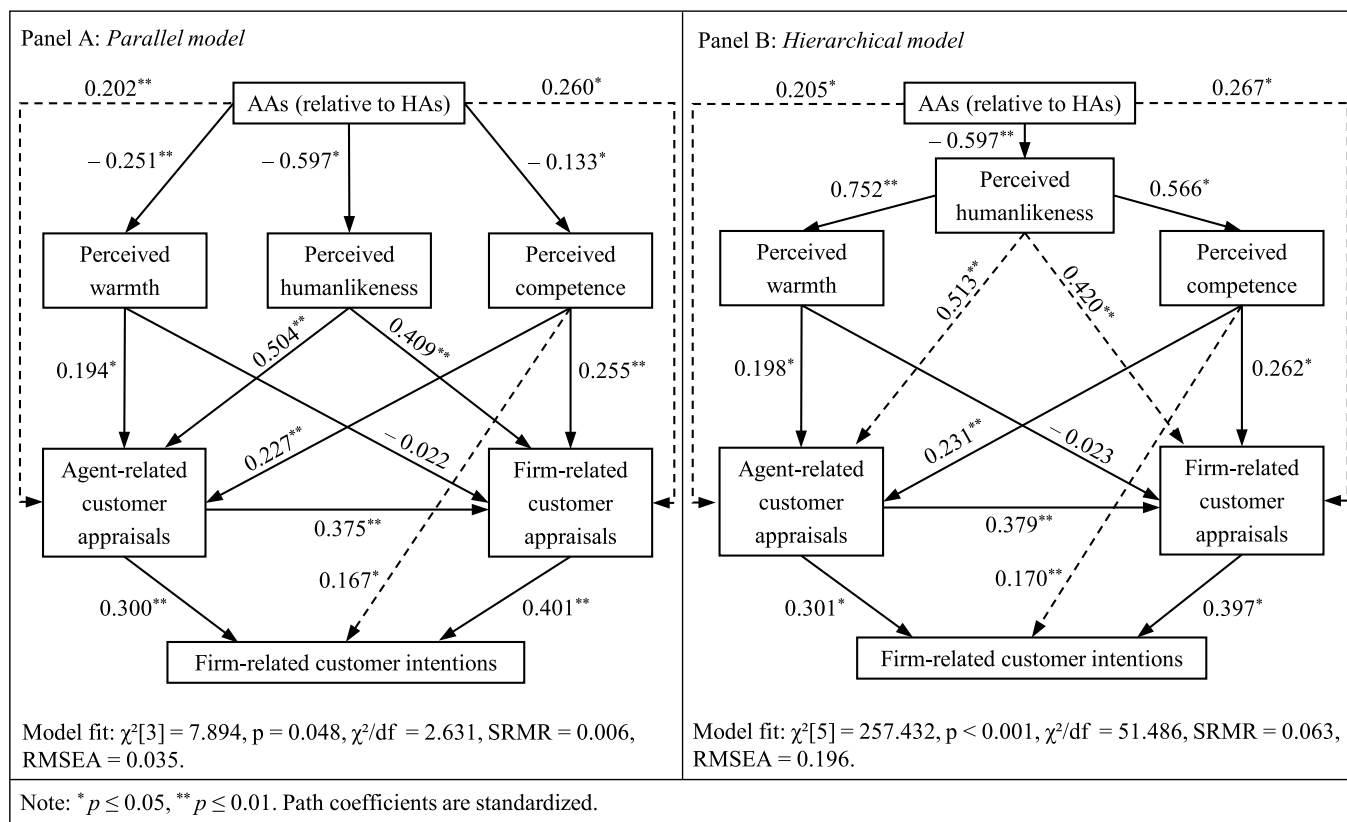


FIGURE 2 | Path model results.

TABLE 3 | Moderation analyses results.

Task-related risks	Functional risk			Psychological risk			Social risk		
	No	Yes	$\Delta\chi^2$	No	Yes	$\Delta\chi^2$	No	Yes	$\Delta\chi^2$
<i>AAs (vs. HAs) →</i>									
• Perceived humanlikeness	-0.603	-0.376	43.903***	-0.641	-0.483	23.092***	-0.555	-0.542	0.148
• Perceived warmth	-0.233	-0.258	0.430	-0.266	-0.189	3.957*	-0.199	-0.375	20.836***
• Perceived competence	-0.167	-0.057	7.927**	-0.163	0.026	22.803***	-0.141	-0.418	52.514***

Note: Standardized path coefficients. Subgroup analyses with χ^2 -difference tests to compare the model fit between an unconstrained and a constrained model (equal path coefficients across subgroups). Significant differences show a worse model fit of the constrained model, reflecting significant differences across subgroups. We additionally tested our hypotheses using Benjamini-Hochberg corrections. The results remained stable, indicating robustness to false discovery rates.

* $p \leq 0.05$; ** $p \leq 0.01$; *** $p \leq 0.001$.

Regarding social risk, the negative effect of AAs (vs. HAs) on perceived warmth is stronger ($p < 0.001$) for tasks with ($\beta = -0.375$, $p = 0.010$) versus without this risk ($\beta = -0.199$, $p = 0.020$). Likewise, the negative effect on perceived competence is stronger ($p < 0.001$) for tasks with ($\beta = -0.418$, $p = 0.005$) versus without this risk ($\beta = -0.141$, $p = 0.032$). These results support H4b and H4c. The result for perceived humanlikeness is nonsignificant ($p = 0.700$), so H4a is not supported.

5 | Conclusion

5.1 | Discussion and Theoretical Implications

5.1.1 | Rethinking Social Perception and Anthropomorphization

Path model results suggest that humanlikeness is not necessarily a prerequisite for warmth and competence. All three

perceptions may emerge in parallel, as a direct response to the AAs versus HAs stimulus. Since perceived humanlikeness was often used as a manipulation check in single studies and rarely linked to other constructs (Table 2), our results provide initial evidence that warrants further robustness tests. Still, our results may prompt a reconsideration of social perception and anthropomorphization. Prior research emphasizes the central role of anthropomorphization in AA-related social perceptions (van Doorn et al. 2017), presenting perceived humanlikeness as a key antecedent of downstream responses (Blut et al. 2021).

The theory of mind perception helps to reassess the assumed sequence “perceived humanlikeness → social perceptions.” It acknowledges that people may anthropomorphize self-propelled nonhuman agents, but describes these agents' interdependence, power, and surveillance (rather than humanlikeness) as conditions that prompt inferences about their concrete intentions and abilities (Epley and Waytz 2010). These

conditions are inherent to AAs, especially those imbued with artificial intelligence (AI). Firms deploy them as default options for customer contact, acting as gatekeepers that autonomously decide whether to forward customers to HAs. Hence, customers need to infer AAs' intentions and abilities. As such, warmth and competence may no longer be perceived as human traits alone, but as assessments of AAs' capacities and use of power.

The interactive abilities and power of AI-enabled AAs may even increase further, as they are trained to read, understand, and anticipate customers' emotions and to react to customers' needs in a personalized way (e.g., Mele et al. 2025). Generative AI, which is based on large language models, is particularly capable in this regard (Huang and Rust 2024). Agentic AI will even enable AAs to execute tasks without step-by-step human instructions, allowing for complex service delivery (Wirtz and Stock-Homburg 2025).

Our results may also nuance speciesism theory. Speciesism describes prejudice or discrimination against nonhuman species (Lafollette and Shanks 1996) and is considered a barrier to AA usage (Schmitt 2020). Recent research suggests that speciesism towards AAs is mitigated when AAs' nonhuman nature is disclosed prior to interactions (De Cicco et al. 2025), when tasks are perceived as beneath humans (Modliński and Trump 2025), and when AAs have low intelligence (Fiestas Lopez Guido et al. 2024). Although speciesism was not measured directly, our results may refine this theory. First, in line with speciesism theory, negative responses to AAs may be primarily driven by their lack of humanlikeness, as this perception has the strongest effect on customer appraisals. However, warmth and competence also exert significant effects (parallel model; Figure 2), suggesting that negative responses to AAs may be to some degree mitigated by signaling warmth (e.g., compliant responses to annoying requests) or competence (e.g., accurate forecasts through superior data-processing abilities). Second, negative responses to AAs depends on task-related risk. Negative judgments are mitigated when customers fear potential underperformance, psychological discomfort, or when bystanders are absent.

5.1.2 | Uncovering the Role of Task-Related Risk

Moderation analyses results show that task-related functional and psychological risk largely weaken the negative effect of AAs (vs. HAs) on agent-related perceptions; social risk largely strengthens it (Table 3). These findings help reconcile diverging views in AA adoption and social psychological research regarding the role of perceived risk in perceptions of AAs. Meta-analyses on AA adoption suggest that AA usage is generally perceived as risky, reducing acceptance (e.g., Mehta et al. 2022). Social psychological theories, while acknowledging that AAs entail uncertainty, posit that people may reduce it by anthropomorphizing such non-human entities and by assessing their warmth and competence, as these inferences help them to familiarize with AAs and their behavior (Epley et al. 2007; Epley and Waytz 2010). This perspective allows for more positive perceptions of AAs. Our results suggest that agent-related perceptions depend on the type of task-related risk, which may work both for and against AAs.

Functional risk works for AAs, weakening their inferiority in perceived humanlikeness and competence; AAs' perceived

competence is almost on par with HAs. Their ability to produce stable outputs may promote positive disconfirmation in task execution. The competence finding partly supports evidence that AAs clearly superior in a task may be evaluated even more positively than HAs (Castelo et al. 2023). The heightened humanlikeness finding is novel: AAs performing a task appropriately (i.e., positive disconfirmation) may itself serve as an anthropomorphization cue, accounting for the weaker negative relationship. For perceived warmth, we observe no moderation effect. Because performance-related tasks may also require warmth (e.g., delivering bad health news), warmth may matter as much as the outcome itself (Yun et al. 2021). Hence, although functional risk favors AAs, customers may still prefer HAs when warmth is essential.

Psychological risk improves AAs' agent-related perceptions. Thus, AAs benefit from tasks that threaten customers' self-esteem, which may seem surprising because they are considered to lack empathy (Castelo et al. 2019). Focusing on task-related psychological risk suggests that customers value AAs' nonjudgmental nature—a beneficial human trait—and therefore perceive AAs as more humanlike and warmer. This insight enriches prior findings that AAs benefit from embarrassing situations because customers experience more positive emotions (de Oliveira Santini et al. 2025). Moreover, AAs' inability to judge even brings them on par with HAs regarding perceived competence, suggesting that customers value AAs' professional handling of inquiries. Collectively, these insights move beyond prior research that links humanoid robot usage (vs. self-service technology) to increased psychological risk (Yoganathan et al. 2021), by showing that task-related psychological risk may favor AAs.

Only social risk works against AAs, amplifying their lower perceived warmth and competence relative to HAs. The presence of bystanders heightens customers' arousal level, resulting in suboptimal task performance that may be attributed to AAs (i.e., self-protection bias). In this sense, customers may blame AAs for being insensitive to their heightened arousal, reducing perceived warmth. Bystanders also impair perceived competence, as AAs' limited cognitive flexibility (Longoni et al. 2019) hinders appropriate responses to spontaneous requests. Contrary to expectations, AAs' lower perceived humanlikeness does not worsen in socially risky tasks, possibly due to a floor effect, as AAs are already perceived as substantially less humanlike than HAs. These findings show the value of examining task-related social risk as a separate moderator. Prior research on AAs considered social risk only as part of an overall risk construct (e.g., Belanche et al. 2025), limiting insights into social risk itself.

5.2 | Managerial Implications

Our findings offer implications for service managers. First, they can increase AAs' perceived humanlikeness, warmth, and competence independently. Prompting humanlikeness, specifically via AAs' appearance (e.g., realistic hair), incurs high design costs, so firms may instead focus on warmth and competence. Regarding warmth, AAs' positive intentions, particularly their nonjudgmental character, could be fostered by, for

example, using non-blameful language. Developers could mitigate discrimination arising from biased training data by using data from multiple cultures or allowing direct employee-customer training (Wirtz and Stock-Homburg 2025). Regarding competence, managers could leverage AAs' data-processing and storage abilities, making them more knowledgeable than HAs. AAs could be trained to recall customers' needs from prior interactions and address them proactively (e.g., re-offering an airport shuttle).

Managers should mind that warmth and competence perceptions may vary across cultures and AI transparency regulations. Most studies came from Western contexts and many also from Eastern contexts, suggesting some generalizability that remains tentative pending future research. Furthermore, regions such as the European Union have specific AI transparency regulations (e.g., disclosing when customers interact with AI-enabled AAs; European Commission 2025). Clearly signaling AAs' machine nature may increase perceived competence regardless of task-related risks, while potentially reducing warmth.

Task-related risks further guide substitution decisions. For functionally risky tasks, firms can leverage AAs' ability to provide stable outcomes, especially when competence is important. For example, firms can communicate AAs' data analysis and forecasting abilities for financial advice. However, when warmth is important for functionally risky tasks (e.g., explaining bad health news) firms should rely on HAs to avoid a loss of human touch. For psychologically risky tasks, firms can substitute HAs with AAs. An illustrative example is the use of AAs as private online tutors. Here, they are not perceived as less competent than HAs. For socially risky tasks, firms should avoid substituting HAs with AAs. Bystanders may cause customers to interact with AAs poorly, whereas HAs can better handle spontaneous questions (e.g., about dishes).

Some tasks may include more than one risk type, where our results suggest opposing recommendations. This is the case when functional or psychological risk occurs along with social risk. As an illustrative example, AAs can advise customers on purchasing a t-shirt but should only give appearance-related feedback when no bystanders are present. Deployment strategies should aim at removing the social risk component.

From an ethical perspective, managerial measures should be implemented with caution. First, since humanlikeness is not necessarily a precondition for AAs to be effective, managers may feel inclined to always replace humans. Second, regarding psychologically risky tasks, AAs are perceived as competent as HAs. Still, in sensitive contexts (e.g., cancer diagnosis or relationship counseling) customers might prefer HAs who can give more profound explanations and advice. Thus, replacing HAs with AAs should be decided on a case-by-case basis, and HAs should remain available alongside the default AA option.

5.3 | Limitations and Future Research

Limitations of this research offer future research opportunities. First, we examined AAs versus HAs. Future research should focus on AAs alone and examine different degrees of agent-related

perceptions. Regarding perceived humanlikeness, an uncanny valley has been suggested. Similarly, optimal degrees of perceived warmth (i.e., when warmth turns into sycophancy) and competence (i.e., when reliability turns into over-precision) may exist.

Second, we did not examine combinations of agent-related perceptions. Future research could examine the combination of different levels of humanlikeness, warmth, and competence for specific tasks. The most favorable combination may be task-specific because the associated expectations for agents may differ. For example, when handing out drugs in a pharmacy (delivering products), AAs might need high (medium) competence, high (low) warmth, and medium (low) humanlikeness.

Third, like any meta-analysis, we used consolidated dependent variables. For example, firm-related customer intentions include variables such as patronage intention. Other customer responses have been omitted because they were rarely addressed (e.g., ethical concerns). Future research could focus on such variables and possibly reveal different result patterns.

Fourth, culture was not examined. Customers' cultural background may impact the extent to which they perceive AAs as humanlike, warm, and competent. This could be due to collectivism. People in Eastern countries, for example, are rather collectivistic (Hofstede 2026) and focus on interpersonal relationships. Thus, they may perceive AAs more negatively than people in Western countries. Future research could examine such moderating effects.

Acknowledgments

Open Access funding enabled and organized by Projekt DEAL.

Funding

The authors have nothing to report.

Ethics Statement

We adhere to the aims and scope of Psychology & Marketing.

Data Availability Statement

The data that support the findings of this study are available from the corresponding author upon reasonable request.

References

- Ajzen, I. 1991. "The Theory of Planned Behavior." *Organizational Behavior and Human Decision Processes* 50, no. 2: 179–211.
- Baker, M. A., J. T. Shin, and Y. W. Kim. 2016. "An Exploration and Investigation of Edible Insect Consumption: The Impacts of Image and Description on Risk Perceptions and Purchase Intent." *Psychology & Marketing* 33, no. 2: 94–112.
- Belanche, D., L. V. Casaló, J. Schepers, and C. Flavián. 2021. "Examining the Effects of Robots' Physical Appearance, Warmth, and Competence in Frontline Services: The Humanness-Value-Loyalty Model." *Psychology & Marketing* 38, no. 12: 2357–2376.
- Belanche, D., L. V. Casaló, M. Flavián, and S. M. C. Loureiro. 2025. "Benefit Versus Risk: A Behavioral Model for Using Robo-Advisors." *Service Industries Journal* 45, no. 1: 132–159.

- Bitner, M. J., B. H. Booms, and L. A. Mohr. 1994. "Critical Service Encounters: The Employee's Viewpoint." *Journal of Marketing* 58, no. 4: 95–106.
- Blut, M., C. Wang, N. V. Wunderlich, and C. Brock. 2021. "Understanding Anthropomorphism in Service Provision: A Meta-Analysis of Physical Robots, Chatbots, and Other AI." *Journal of the Academy of Marketing Science* 49, no. 4: 632–658.
- Bond, Jr., C. F., and L. J. Titus. 1983. "Social Facilitation: A Meta-Analysis of 241 Studies." *Psychological Bulletin* 94, no. 2: 265–292.
- Castelo, N., M. W. Bos, and D. R. Lehmann. 2019. "Task-Dependent Algorithm Aversion." *Journal of Marketing Research* 56, no. 5: 809–825.
- Castelo, N., J. Boegershausen, C. Hildebrand, and A. P. Henkel. 2023. "Understanding and Improving Consumer Reactions to Service Bots." *Journal of Consumer Research* 50, no. 4: 848–863.
- Clegg, M., R. Hofstetter, E. de Bellis, and B. H. Schmitt. 2024. "Unveiling the Mind of the Machine." *Journal of Consumer Research* 51, no. 2: 342–361.
- De Cicco, R., M. G. Elmashhara, S. C. Silva, and M. Hammerschmidt. 2025. "The Impact of Providing Non-Human Identity Cues About Sales Agents on Consumer Responses: The Role of Social Presence and Speciesism Activation." *European Journal of Marketing* 59, no. 13: 55–84.
- Derbaix, C. 1983. "Perceived Risk and Risk Relievers: An Empirical Investigation." *Journal of Economic Psychology* 3, no. 1: 19–38.
- van Doorn, J., M. Mende, S. M. Noble, et al. 2017. "Domo Arigato Mr. Roboto: Emergence of Automated Social Presence in Organizational Frontlines and Customers' Service Experiences." *Journal of Service Research* 20, no. 1: 43–58.
- van Doorn, J., G. Odekerken-Schröder, and J. Spohrer. 2025. "Robots Are Here to Stay: Time to Invest in a Future We Actually Want to Live in." *Journal of Service Research* 28, no. 1: 3–8.
- Epley, N., and A. Waytz. 2010. "Mind Perception." *Handbook of Social Psychology* 1, no. 5: 498–541.
- Epley, N., A. Waytz, and J. T. Cacioppo. 2007. "On Seeing Human: A Three-Factor Theory of Anthropomorphism." *Psychological Review* 114, no. 4: 864–886.
- European Commission. 2025. "Guidelines and Code of Practice on Transparent AI Systems," (accessed March 27, 2026). <https://digital-strategy.ec.europa.eu/en/faqs/guidelines-and-code-of-practice-transparent-ai-systems>.
- Fiestas Lopez Guido, J. C., J. W. Kim, P. T. L. Popkowski Leszczyc, N. Pontes, and S. Tuzovic. 2024. "Retail Robots as Sales Assistants: How Speciesism Moderates the Effect of Robot Intelligence on Customer Perceptions and Behaviour." *Journal of Service Theory and Practice* 34, no. 1: 127–154.
- Fiske, S. T., A. J. C. Cuddy, and P. Glick. 2007. "Universal Dimensions of Social Cognition: Warmth and Competence." *Trends in Cognitive Sciences* 11, no. 2: 77–83.
- Fornell, C. 1983. "Issues in the Application of Covariance Structure Analysis: A Comment." *Journal of Consumer Research* 9, no. 4: 443–448.
- Gelbrich, K., H. Roschk, S. Miederer, and A. Kerath. 2026. "Automated Versus Human Agents: A Meta-Analysis of Customer Responses to Robots, Chatbots, and Algorithms and Their Contingencies." *Journal of Marketing* 90, no. 2: 1–26.
- Gray, H. M., K. Gray, and D. M. Wegner. 2007. "Dimensions of Mind Perception." *Science (New York, N.Y.)* 315, no. 5812: 619.
- Heerink, M., B. Kröse, V. Evers, and B. Wielinga. 2008. "The Influence of Social Presence on Acceptance of a Companion Robot by Older People." *Journal of Physical Agents (JoPha)* 2, no. 2: 33–40.
- Heider, F. 1958. *The Psychology of Interpersonal Relations*. New York: John Wiley & Sons.
- Hofstede, G. J. 2026. "Country Comparison Bar Charts," (accessed March 13, 2026). <https://geerthofstede.com/country-comparison-bar-charts/>.
- Holthöwer, J., and J. van Doorn. 2023. "Robots Do Not Judge: Service Robots Can Alleviate Embarrassment in Service Encounters." *Journal of the Academy of Marketing Science* 51: 767–784.
- Hoyer, W. D., D. J. MacInnis, and R. Pieters. 2024. *Consumer Behavior* (8th ed.). Boston, MA: Cengage Learning.
- Huang, M. H., and R. T. Rust. 2024. "The Caring Machine: Feeling AI for Customer Care." *Journal of Marketing* 88, no. 5: 1–23.
- Im, H., B. Sung, G. Lee, and K. Q. Xian Kok. 2023. "Let Voice Assistants Sound Like a Machine: Voice and Task Type Effects on Perceived Fluency, Competence, and Consumer Attitude." *Computers in Human Behavior* 145: 107791.
- Imhof, G., and P. Klaus. 2020. "The Dawn of Traditional CX Metrics? Examining Satisfaction, EXQ, and WAR." *International Journal of Market Research* 62, no. 6: 673–688.
- Jöreskog, K. G. 1971. "Simultaneous Factor Analysis in Several Populations." *Psychometrika* 36, no. 4: 409–426.
- Lafollette, H., and N. Shanks. 1996. "The Origin of Speciesism." *Philosophy (London, England)* 71, no. 275: 41–61.
- Landis, J. R., and G. G. Koch. 1977. "The Measurement of Observer Agreement for Categorical Data." *Biometrics* 33, no. 1: 159–174.
- Larkin, C., C. Drummond Otten, and J. Árvai. 2022. "Paging Dr. JAR-VIS! Will People Accept Advice From Artificial Intelligence for Consequential Risk Management Decisions?" *Journal of Risk Research* 25, no. 4: 407–422.
- Liadeli, G., F. Sotgiu, and P. W. J. Verlegh. 2023. "A Meta-Analysis of the Effects of Brands' Owned Social Media on Social Media Engagement and Sales." *Journal of Marketing* 87, no. 3: 406–427.
- Lipsey, M. W., and D. B. Wilson. 2001. *Practical Meta-Analysis*. SAGE Publications.
- Longoni, C., A. Bonezzi, and C. K. Morewedge. 2019. "Resistance to Medical Artificial Intelligence." *Journal of Consumer Research* 46, no. 4: 629–650.
- Mehta, P., C. Jebarajakirthy, H. I. Maseeh, A. Anubha, R. Saha, and K. Dhanda. 2022. "Artificial Intelligence in Marketing: A Meta-Analytic Review." *Psychology & Marketing* 39, no. 11: 2013–2038.
- Mele, C., T. Russo-Spena, I. Di Bernardo, and S. Gherardi. 2025. "Affect-Based Well-Being in Caring Practices With Companion Robots." *Journal of Service Research*: 1–18.
- Modliński, A., and R. K. Trump. 2025. "The Impact of Speciesism on Customers' Acceptance of Service Automation." *Journal of Service Theory and Practice* 35, no. 2: 245–262.
- Murray, K. B., and J. L. Schlacter. 1990. "The Impact of Services Versus Goods on Consumers' Assessment of Perceived Risk and Variability." *Journal of the Academy of Marketing Science* 18, no. 1: 51–65.
- Nicol, A. A. M., and K. de France. 2024. "Nonjudgmental Regard of Others: Investigating the Links Between Other-Directed Trait Mindfulness and Prejudice." *Psychological Reports* 127, no. 1: 178–211.
- de Oliveira Santini, F., W. M. Lim, C. H. Sampaio, et al. 2025. "The Robot-Human Paradox: A Meta-Analysis of Customer Service by Robots Versus Humans on Customer Experience." *Journal of Consumer Behaviour* 24, no. 3: 1392–1404.
- Peterson, R. A., and S. P. Brown. 2005. "On the Use of Beta Coefficients in Meta-Analysis." *Journal of Applied Psychology* 90, no. 1: 175–181.
- Pitardi, V., J. Wirtz, S. Paluch, and W. H. Kunz. 2022. "Service Robots, Agency, and Embarrassing Service Encounters." *Journal of Service Management* 33, no. 2: 389–414.

- Pyszczynski, T., J. Greenberg, S. Solomon, J. Arndt, and J. Schimel. 2004. "Why Do People Need Self-Esteem? A Theoretical and Empirical Review." *Psychological Bulletin* 130, no. 3: 435–468.
- Rubera, G., and A. H. Kirca. 2012. "Firm Innovativeness and Its Performance Outcomes: A Meta-Analytic Review and Theoretical Integration." *Journal of Marketing* 76, no. 3: 130–147.
- Schmitt, B. 2020. "Speciesism: An Obstacle to AI and Robot Adoption." *Marketing Letters* 31, no. 1: 3–6.
- Sheng, M. L., N. Natalia, and E. Z. Rusfian. 2025. "AI Chatbot, Human, and In-Between: Examining the Broader Spectrum of Technology-Human Interactions in Driving Customer-Brand Relationships Across Experience and Credence Services." *Psychology & Marketing* 42, no. 4: 1051–1071.
- Sohn, S. 2024. "Consumer Perceived Risk of Using Autonomous Retail Technology." *Journal of Business Research* 171: 114389.
- Srinivasan, R., and G. Sarial-Abi. 2021. "When Algorithms Fail: Consumers' Responses to Brand Harm Crises Caused by Algorithm Errors." *Journal of Marketing* 85, no. 5: 74–91.
- Taylor, J. W., and J. W. Taylor. 1974. "The Role of Risk in Consumer Behavior: A Comprehensive and Operational Theory of Risk Taking in Consumer Behavior." *Journal of Marketing* 38, no. 2: 54–60.
- Viechtbauer, W. 2010. "Conducting Meta-Analyses in R With the Metafor Package." *Journal of Statistical Software* 36, no. 3: 1–48.
- Wentzel, D. 2009. "The Effect of Employee Behavior on Brand Personality Impressions and Brand Attitudes." *Journal of the Academy of Marketing Science* 37, no. 3: 359–374.
- Wirtz, J., and R. Stock-Homburg. 2025. "Generative AI Meets Service Robots." *Journal of Service Research* 28: 527–543.
- Wirtz, J., P. G. Patterson, W. H. Kunz, et al. 2018. "Brave New World: Service Robots in the Frontline." *Journal of Service Management* 29, no. 5: 907–931.
- Yoganathan, V., V. S. Osburg, W. H. Kunz, and W. Toporowski. 2021. "Check-in at the Robo-Desk: Effects of Automated Social Presence on Social Cognition and Service Implications." *Tourism Management* 85: 104309.
- Yun, J. H., E.-J. Lee, and D. H. Kim. 2021. "Behavioral and Neural Evidence on Consumer Responses to Human Doctors and Medical Artificial Intelligence." *Psychology & Marketing* 38, no. 4: 610–625.
- Zajonc, R. B., and S. M. Sales. 1966. "Social Facilitation of Dominant and Subordinate Responses." *Journal of Experimental Social Psychology* 2, no. 2: 160–168.
- Zehnle, M., C. Hildebrand, and A. Valenzuela. 2025. "Not All AI Is Created Equal: A Meta-Analysis Revealing Drivers of AI Resistance Across Markets, Methods, and Time." *International Journal of Research in Marketing* 42, no. 3: 729–751.
- Zeithaml, V. A., M. J. Bitner, D. D. Gremler, and M. Mende. 2024. "Services Marketing: Integrating Customer Focus Across the Firm." In *International Edition* (8th ed.). New York, NY: McGraw-Hill.

Supporting Information

Additional supporting information can be found online in the Supporting Information section.

Supporting File: mar70216-sup-0001-Web_Appendix.docx.