

Research note

From cafés to clinics: Consumer attitudes toward human-like and machine-like service robot failures[☆]

Ezgi Merdin-Uygur^{a,1} , Selcen Ozturkcan^{b,c,*}

^a Brunel Business School, Brunel University of London, Eastern Gateway, Kingston Lane, Uxbridge UB8 3PN, UK

^b School of Business and Economics, Linnaeus University, Sweden

^c Sabanci Business School, Sabanci University, Turkey

ARTICLE INFO

Keywords:

Robotic service agents
Chatbot
Service robot
Hospitality services
Anthropomorphism
Robotic failure

ABSTRACT

This study examines consumer evaluations of robotic service failures caused by human interference by integrating service context, robot appearance, and individual anthropomorphism tendencies into a unified model. Two between-subjects experiments were conducted. In Study 1 ($N = 402$), participants interacted with a healthcare or food-service bot that failed due to verbal interference. Healthcare service failure elicited significantly more negative attitudes and lower failure tolerance than food service failure, and failure tolerance fully mediated the relationship between context and attitudes. In Study 2 ($N = 213$), we employed a 2×2 design (healthcare vs. food services $\times$ human-like vs. machine-like robot) and measured perceived deservingness and trait anthropomorphism. Human-like robots were judged most harshly when failing in healthcare (vs. food) services, whereas machine-like robots received similar evaluations across contexts. Perceived deservingness of the robot mediated this interaction. Moreover, the moderated-mediation effect occurred only among individuals with low to medium anthropomorphism tendencies. By positioning failure tolerance and deservingness judgments as core mechanisms in human-robot interaction, our findings advance theoretical understanding of moral attributions in service failure. Practically, they highlight the importance of matching robot anthropomorphic cues to service criticality: less human-like designs in high-stakes environments, while more human-like appearances may be appropriate in lower-stakes settings.

1. Introduction

As robotic service agents in hospitality and healthcare industries increase, their failures also become unavoidable. Physical human interventions to make agents fail are underexplored compared to system errors (Lteif and Valenzuela, 2022). We shift the perspective from robots harming humans (Swiderska and Küster, 2020) to human abuse. Critical variables in such failures, like service context and anthropomorphism, were examined in isolation (i.e., Huang et al., 2024). We examine the interaction between service context and robot appearance on post-failure consumer evaluations. Table 1 encapsulates previous research and our unique contribution.

Limited research shows that service type shapes consumer responses to failures (Liu et al., 2022). For example, private vs. public service

context moderates consumer tolerance for robotic failures (Ma et al., 2024). In healthcare, repercussions and expectations rise (Weun et al., 2004), and tolerance drops, unlike in restaurants, where minor errors are tolerated. We hypothesise:

H1. *Robotic agents providing healthcare (vs. food) services elicit more negative post-failure attitudes.*

H2. *Service failure tolerance mediates the relationship between service type and post-failure attitudes.*

Consumers assess robotic service failures based on perceived human-likeness (Choi et al., 2021). Humanoid robots raise expectations, causing greater dissatisfaction when they fail (Wirtz et al., 2018). Conversely, machine-like robots encounter diminished emotional expectations and

[☆] No conflicting interests.

* Correspondence to: School of Business and Economics, Linnaeus University, 352 52, Växjö, Sweden.

E-mail addresses: ezgi.merdinuygur@brunel.ac.uk (E. Merdin-Uygur), Selcen.Ozturkcan@lnu.se, selcen@sabanciuniv.edu (S. Ozturkcan).

¹ <https://orcid.org/0000-0002-4065-7336>.

² <https://orcid.org/0000-0003-2248-0802>.

undergo less scrutiny (Zhang et al., 2023). In high-stakes contexts such as healthcare, emotional connections may be inappropriate or even detrimental (Weun et al., 2004). Yet, in low-stakes environments such as cafés, demand is minimal and consumers tend to be more understanding of machine-like agents (Wirtz et al., 2018).

H3. *The anthropomorphic cues of the robotic service agent moderate the relationship between service type and post-failure attitudes in such a way that the service type influences post-failure attitudes only for human-like agents.*

High anthropomorphisers are more inclined to ascribe human-like characteristics to machines (Zheng et al., 2025) and tend to forgive failures of humanoids (Murray, 2022). We hypothesise the interactive role of robot anthropomorphism with consumers' trait anthropomorphism as:

H4. *Consumers' trait anthropomorphism moderates the anthropomorphism and service type interaction on post-failure attitudes in such a way that the interaction influences post-failure attitudes only for those with high trait anthropomorphism.*

Deservingness judgments are pivotal in assessing robotic success or failure. In line with Just World Theory, individuals believe that good or bad outcomes are deserved (Callan et al., 2014). Deservingness denotes the belief that an individual is entitled to a specific outcome (Palmeira et al., 2022), and anthropomorphic designs amplify such reactions (Ward et al., 2013).

H5. *Perceived deservingness mediates the effect of the interaction between anthropomorphism and service type on post-failure attitudes.*

Fig. 1 encapsulates our proposed model.

2. Method

We employed two between-subjects experimental designs manipulating service context (Study 1, 2) and anthropomorphism (Study 2). Table 2 summarises our measures. 402 (267 women; $M_{\text{age}} = 38.06$) participants for Study 1 and 213 participants (103 women; $M_{\text{age}} = 41.70$) for Study 2 were recruited via Prolific. Most reported middle-income ($\approx 40\%$) and held a bachelor's degree ($\approx 46\text{--}50\%$).

Healthcare and food-service contexts also reflect contemporary trends within robotic hospitality services (Assiouras et al., 2025).

Study 1 and Results: Participants were randomly assigned to *healthcare* or *food service* chatbots that failed due to verbal human interference. Chatbot screenshots were adapted from Pavone et al. (2023), demonstrating the dialogue:

'ALEX' THE BOT ASSISTANT: How can I assist you today regarding your food menu (vs. medicine) choices?

'CUSTOMER 9847387': You are dumb and don't know anything about real food (vs. medicine)!

'ALEX' THE BOT ASSISTANT: ...Error: Conversation terminated. Please contact support.

Measures included manipulation checks, attitudes, and failure tolerance (Table 2). All participants correctly identified the bot type and failure and found scenarios equally credible ($p = .27$). As hypothesised, participants showed less favourable attitudes ($M = 4.15$ vs. 4.51 ; $t(1,400) = -2.48$, $p = .01$) and lower tolerance ($M = 3.78$ vs. 4.00 ; $t(1,400) = -2.13$, $p = .03$) toward the healthcare (vs. food) bot. PROCESS Macro Model 4 (Preacher and Hayes, 2008) confirmed mediation by service failure tolerance ($B = .18$, $SE = .09$, 95% CI $[0.02, 0.35]$). We eliminated confounding explanation of service criticality ($M = 4.29$ vs. 4.11 ; $t(1,400) = -1.10$, $p = .27$).

Study 2 and Results: Participants were randomly assigned to one of four conditions (human-like vs. machine-like robot; serving medicine vs. food). They viewed a robot illustration dropping food vs. medicine after human interference, completed manipulation checks, attitude, deservingness, and the trait-anthropomorphism scales.

204/213 participants correctly identified the scenario as a failure ($\chi^2(1) = 192.93$, $p < .01$). All participants recalled the healthcare setting correctly, while 128/133 in the cafe condition (96.2%) did so. All scenarios were rated equally credible ($p = .83$).

Robot appearance and service context interacted to affect perceived deservingness ($F(1,209) = 4.91$, $p = .03$) and attitudes $F(1,209) = 2.97$, $p = .09$, marginal significance). Human-like robots were perceived as less deserving of failure in cafés vs. healthcare services ($M = 2.21$ vs. $M = 2.74$, $SE = .20$, $p = .01$) with no difference for machine-like robots ($M = 2.43$ vs. 2.32 , $SE = .20$, $p = .61$). Human-like robots were evaluated more positively when they failed in a café vs. healthcare services ($M = 4.43$ vs. $M = 3.55$; $SE = .30$, $p = .00$), with no difference for machine-like robots ($M = 4.25$ vs. 4.09 ; $SE = .30$, $p = .59$).

PROCESS Macro Model 9 (Preacher and Hayes, 2008) confirmed the full moderated mediation model ($B = .36$, $SE = .15$, 95% CI $[0.07, 0.66]$). Conditional effects were significant for low ($B = .93$, $SE = .24$, 95% CI $[0.45, 1.40]$) or medium anthropomorphisers ($B = .20$, $SE = 3.00$, 95% CI $[0.21, 0.99]$).

3. Discussion and conclusions

In conclusion, human-like robots were evaluated more critically when they failed in healthcare environments compared to food settings, whereas this phenomenon did not extend to machine-like robots. Our key findings are presented in Table 3.

Our research shows that in low-tolerance settings (e.g., healthcare), humanoid robot failures intensify negative evaluations, reinforcing pressures facing anthropomorphised agents (Fan et al., 2020). We further demonstrate that perceived deservingness acts as a moral filter through which consumers react to failure, meaningful for future human-robot interaction research and expanding Liu et al. (2023) social framework. We contribute robotic service recovery and consumer forgiveness frameworks (Nguyen et al., 2025). Although

Table 1
Select works on robotic services.

Ref.	Service Type	Service Failure	Anthropomorphism	Theoretical Mechanism
Cheng, 2023	-	+	+	internal attribution
Choi et al., 2021	-	+	+	perceived warmth
Fan et al., 2020	-	+	+	self-blame
Huang et al., 2024	female- vs. male-dominated	+	-	communion and tolerance
Liu et al., 2022	hedonic vs. utilitarian	-	child-like vs. adult-like	warmth/competence and trust
Ma et al., 2024	instrumental vs. assistive vs. companionate role	+	+	tolerance and psychological distance
Merdin-Uygur and Ozturkcan, 2025	-	+	+	deservingness, self-efficacy regarding robots
Seo, 2022	-	-	+	pleasure
Sheng et al., 2024	-	+	+	empathy
This paper	+	+	+	deservingness and tolerance

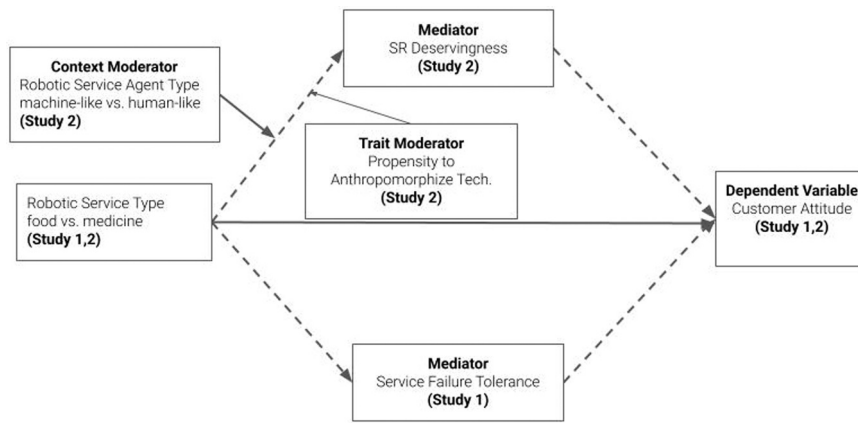


Fig. 1. Proposed research model.

Table 2
Summary of measures.

Study	Construct	Number of items	Response formats	Scale	Reliability (Cronbach's Alpha)	
					Study 1	Study 2
1, 2	Attitudes (Hesapci et al., 2016)	Six-item	"Please rate how you feel about this robot in terms of these dimensions"	7-point semantic differential scale: irritating/not irritating; not appealing/appealing; unlikeable/likeable; bad/good; negative/positive; unfavourable/favourable.	0.95	0.96
1, 2	Involvement	Single-item	"How involved/interested are you with AI/robotics in general?"	7-point Likert scale: 1 = not interested at all; 7 = very much interested.	N/A	
1	Service failure tolerance	Three-item	"I think the robotic agent's service failure was unforgivable." (reverse item) "I would share this negative experience with others." (reverse item) "I would recommend the robotic agent to others."	7-point Likert scale: 1 = strongly disagree; 7 = strongly agree.	0.51*	
1	Service criticality (Mozafari et al., 2022)	Single-item	The resolution of the service request from the robotic agent is...	7-point semantic differential scale: uncritical/ critical	N/A	
2	Deservingness (Callan et al., 2014)	Four-item	"the robot deserves the bad that happens to it" "the robot is deserving of what happened" "the robot deserves to do poorly" "the robot is deserving of positive outcomes" (reverse item)	7-point Likert scale: 1 = strongly disagree; 7 = strongly agree.		0.78
2	Propensity to anthropomorphise technology (Waytz et al., 2010)	Three-item	"To what extent do you believe that...: technology has intentions? an AI chatbot can experience emotions? the average AI chatbot has consciousness?"	5-point Likert scale: 1 = not at all; 7 = to a great extent		0.74

* No item deletion increases reliability.

Table 3
Key findings from study 1 and study 2.

Hypothesis	Finding	Study
H1	Less favourable post-failure attitudes toward healthcare agents	1
H2	Mediator: Tolerance for failure	1
H3	Moderator: Effect valid only for human-like agents	2
H4	Moderator: Effect valid only for low/medium anthropomorphisers	2
H5	Mediator: Deservingness	2

anthropomorphism facilitates forgiveness (Zhao et al., 2025), it is contingent upon context. Robotic designers need to calibrate anthropomorphic cues (i.e., facial features) in high-pressure settings.

CRediT authorship contribution statement

Selcen Ozturkcan: Writing – review & editing, Writing – original draft, Visualization, Validation, Supervision, Methodology, Conceptualization. **Ezgi Merdin-Uygur:** Writing – review & editing, Writing – original draft, Visualization, Validation, Supervision, Software, Methodology, Funding acquisition, Formal analysis, Data curation.

Declaration of Competing Interest

Authors have nothing to declare.

Acknowledgement

The authors would like to note appreciation for the grant received by the first author from the Marketing Trust in support of this research.

Data availability

Data will be made available on request.

References

- Assiouras, I., Laserer, C., Buhalis, D., 2025. The evolution of artificial empathy in the hospitality metaverse era. *Int. J. Hosp. Manag.* 126, 104063.
- Callan, M.J., Kay, A.C., Dawtry, R.J., 2014. Making sense of misfortune: deservingness, self-esteem, and patterns of self-defeat. *J. Personal. Soc. Psychol.* 107 (1), 142–162.
- Cheng, L.K., 2023. Effects of service robots' anthropomorphism on consumers' attribution toward and forgiveness of service failure. *J. Consum. Behav.* 22 (1), 67–81.
- Choi, S., Mattila, A.S., Bolton, L.E., 2021. To err is human(-oid): how do consumers react to robot service failure and recovery? *J. Serv. Res.* 24 (3), 354–371.
- Fan, A., Wu, L., Miao, L., Mattila, A.S., 2020. When does technology anthropomorphism help alleviate customer dissatisfaction after a service failure? – the moderating role of consumer technology self-efficacy and interdependent self-construal. *J. Hosp. Mark. Manag.* 29 (3), 269–290.
- Hesapci, O., Merdin, E., Gorgulu, S., 2016. Your ethnic model speaks to the culturally connected: Differential effects of model ethnicity in advertisements and the role of cultural self-construal. *J. Consum. Behav.* 15 (2), 175–185. <https://doi.org/10.1002/cb.1562>.
- Huang, H., Chen, F.F., Liu, S.Q., 2024. Gender stereotyping in robot service failures. *J. Hosp. Tour. Res.* (ja).
- Liu, D., Li, C., Zhang, J., Huang, W., 2023. Robot service failure and recovery: Literature review and future directions. *Int. J. Adv. Robot. Syst.* 20 (4), 17298806231191606. <https://doi.org/10.1177/17298806231191606>.
- Liu, X.(S.), Yi, X. (S.), Wan, L.C., 2022. Friendly or competent? The effects of perception of robot appearance and service context on usage intention. *Ann. Tour. Res.* 92, 103324.
- Lteif, L., Valenzuela, A., 2022. The effect of anthropomorphised technology failure on the desire to connect with others. *Psychol. Mark.* 39 (9), 1762–1774.
- Ma, K., Duan, X., Fu, X., Liu, W., Zheng, M., 2024. Effects of human-robot interaction type on customer tolerance of humanoid robot service failure. *J. Hosp. Mark. Manag.*
- Merdin-Uygur, E., Ozturkcan, S., 2025. The robot saw it coming: physical human interference, deservingness, and self-efficacy in service robot failures. *Behav. Inf. Technol.*
- Mozafari, N., Weiger, W.H., Hammerschmidt, M., 2022. Trust me, I'm a bot–repercussions of chatbot disclosure in different service frontline settings. *J. Serv. Manag.* 33 (2), 221–245. <https://doi.org/10.1108/josm-10-2020-0380>.
- Murray, R.P., 2022. Self-Service Technologies That Appear Human Interacting with Customers: Effects on Third-Party Observers (Doctoral dissertation). The University of Texas Rio Grande Valley.
- Nguyen, H.N., Nguyen, N.T., Hancer, M., 2025. Human-robot collaboration in service recovery: Examining apology styles, comfort emotions, and customer retention. *Int. J. Hosp. Manag.* 126, 104028.
- Palmeira, M., Koo, M., Sung, H.A., 2022. You deserve the bad (or good) service: the role of moral deservingness in observers' reactions to service failure (or excellence). *Eur. J. Mark.* 56 (3), 653–676.
- Pavone, G., Meyer-Waarden, L., Munzel, A., 2023. Rage against the machine: experimental insights into customers' negative emotional responses, attributions of responsibility, and coping strategies in artificial intelligence–based service failures. *J. Interact. Mark.* 58 (1), 52–71.
- Preacher, K.J., Hayes, A.F., 2008. Contemporary approaches to assessing mediation in communication research. In: Hayes, A.F., Slater, M.D., Snyder, L.B. (Eds.), *The Sage Sourcebook of Advanced Data Analysis Methods for Communication Research*. Sage Publications, Inc, pp. 13–54.
- Seo, S., 2022. When female (male) robot is talking to me: effect of service robots' gender and anthropomorphism on customer satisfaction. *Int. J. Hosp. Manag.* 102, 103166.
- Sheng, X., Murray, R., Ketron, S.C., Felix, R., 2024. Effects of third-party observer empathy when viewing interactions between robots and customers: the moderating role of robot eeriness. *Int. J. Hosp. Manag.* 119, 103729.
- Swiderska, A., Küster, D., 2020. Robots as malevolent moral agents: harmful behavior results in dehumanization, not anthropomorphism. *Cogn. Sci.* 44 (7), e12872 (n/a).
- Ward, A.F., Olsen, A.S., Wegner, D.M., 2013. The harm-made mind: observing victimization augments attribution of minds to vegetative patients, robots, and the dead. *Psychol. Sci.* 24 (8), 1437–1445.
- Waytz, A., Cacioppo, J., Epley, N., 2010. Who sees human? The stability and importance of individual differences in anthropomorphism. *Perspect. Psychol. Sci.* 5 (3), 219–232.
- Weun, S., Beatty, S.E., Jones, M.A., 2004. The impact of service failure severity on service recovery evaluations and post-recovery relationships. *J. Serv. Mark.* 18 (2), 133–146.
- Wirtz, J., Patterson, P.G., Kunz, W.H., Gruber, T., Lu, V.N., Paluch, S., Martins, A., 2018. Brave new world: service robots in the frontline. *J. Serv. Manag.* 29 (5), 907–931.
- Zhang, M., Cui, J., Zhong, J., 2023. How consumers react differently toward humanoid vs. nonhumanoid robots after service failures: a moderated chain mediation model. *Int. J. Emerg. Mark.*
- Zhao, Y., Zhu, Z., Tang, B., 2025. Are highly anthropomorphic service robots more likely to be forgiven by customers after service failures? A mind perception perspective. *Int. J. Hosp. Manag.* 126, 104103.
- Zheng, C., Liu, S., Zhao, L., Ma, K., Wang, W., Wang, H., 2025. How to make consumers tolerate robotic service failures. *Int. J. Hosp. Manag.* 126, 104059.