



Consumer Trust Formation in Artificial Intelligence–Enabled Service Encounters

Prof. Ahmad Farid

Professor, Denesan University, Kuala Lumpur, Malaysia

Abstract

Artificial intelligence has become an increasingly visible component of contemporary service encounters through conversational agents, recommendation systems, automated support interfaces, intelligent service robots, virtual assistants, and decision-support technologies. While these applications can improve responsiveness, personalization, accessibility, and operational efficiency, their effectiveness depends substantially on whether consumers consider the AI-enabled interaction sufficiently trustworthy. Trust formation in such encounters is complex because consumers simultaneously evaluate the functional capability of the AI system, the intentions and accountability of the service provider, the fairness of algorithmic processes, the security of personal information, and the availability of human assistance when automated interaction becomes inadequate.

This study develops an integrated framework for understanding consumer trust formation in artificial intelligence–enabled service encounters. The framework conceptualizes trust as an evolving judgment influenced by perceived competence, responsiveness, appropriate transparency, fairness, privacy assurance, and accessible human support. Anthropomorphic and empathetic characteristics are treated as contextual trust cues rather than universally beneficial design features. Because no original consumer dataset was supplied, the analytical component employs a transparent simulation-based design involving 600 synthetic service-interaction profiles divided equally across low, moderate, and high trust-support conditions. A Consumer Trust Index was modeled on a standardized 0–100 scale. The simulated analysis produced trust scores of 48.5, 67.1, and 82.8 across the three conditions, indicating a coherent theoretical relationship between stronger assurance mechanisms and greater modeled consumer trust. These figures are methodological outputs and should not be interpreted as population estimates or causal effects.

The study further argues that organizations should pursue calibrated trust rather than indiscriminate trust maximization. Transparency without understandable explanation, anthropomorphism without competence, personalization without privacy assurance, or automation without effective human escalation may weaken rather than strengthen consumer confidence. The proposed framework contributes to service and consumer-behavior research by integrating technical, relational, ethical, and organizational trust mechanisms within a single service-encounter perspective and offers propositions for future empirical validation.



Keywords: Consumer Trust, Artificial Intelligence, AI-Enabled Services, Service Encounters, Chatbots, Transparency, Perceived Competence, Algorithmic Fairness, Privacy Assurance, Human–AI Interaction

1. Introduction

Artificial intelligence is increasingly embedded in customer-facing service processes, ranging from conversational customer support and automated recommendations to voice assistants, service robots, financial decision tools, healthcare applications, and intelligent digital agents. Service research has long recognized that AI can undertake increasingly sophisticated combinations of mechanical, analytical, intuitive, and relational service tasks, thereby changing how organizations design and deliver customer experiences. Huang and Rust's influential service framework anticipated this progressive expansion of AI capability and emphasized that different types of intelligence require different combinations of human and machine involvement.

The expansion of AI-mediated service has made consumer trust a central managerial and theoretical issue. An AI-enabled encounter creates a distinctive form of vulnerability because consumers may need to rely on recommendations, disclose personal information, follow automated instructions, accept algorithmic decisions, or purchase products without fully understanding how the system reached its conclusion. Reviews of human trust in artificial intelligence identify reliability, transparency, system capability, representation, and interaction characteristics as recurring determinants of trust. More recent customer-service research similarly distinguishes functional or cognitive trust—associated with accuracy, competence, responsiveness, security, and transparency—from relational or affective trust associated with naturalness, empathy, personalization, and social presence.

Trust formation, however, cannot be reduced to making an AI system appear more human. Anthropomorphic characteristics may increase social presence and make conversational interaction easier, yet their influence varies according to consumer expectations, relationship norms, service context, perceived competence, and the degree of human-likeness. Research examining anthropomorphism in chatbot interaction has found that human-like characteristics can influence trust, although the relationship is conditional rather than universally positive. More recent evidence likewise suggests that near-human or highly human-like interfaces may not always generate the strongest trust response.

Transparency produces a similarly nuanced relationship. Consumers often expect organizations to disclose when AI is involved and provide sufficient information about automated processes, yet disclosure alone does not automatically generate trust. Luo and colleagues' field research demonstrated that chatbot identity disclosure could reduce purchase behavior when consumers interpreted the AI as less knowledgeable and less empathetic. More recent experimental work shows that algorithmic transparency can function as a signal of organizational accountability, but its influence on trust in the AI system itself may depend on pre-existing attitudes and issue involvement. Other experimental research has reported a broader “transparency dilemma,” in which disclosing AI use can sometimes reduce trust rather than increase it.

These findings expose an important research gap. Many studies examine isolated antecedents such as anthropomorphism, transparency, privacy, service quality, or purchase intention, whereas consumers experience AI-enabled services as integrated encounters in which technical performance,



ethical assurance, interpersonal design, organizational reputation, and recovery mechanisms operate simultaneously. The present study therefore develops an integrated model of Consumer Trust Formation in Artificial Intelligence–Enabled Service Encounters. It proposes that sustainable trust develops when consumers perceive the AI as competent and responsive, consider the process sufficiently transparent and fair, believe that personal information is appropriately protected, and retain practical access to accountable human intervention. The study deliberately distinguishes trustworthy service design from simple trust maximization and uses simulation-based analysis to demonstrate the internal behavior of the proposed framework without presenting synthetic observations as empirical consumer data.

2. Objectives of the Study

2.1 To Identify the Principal Drivers of Consumer Trust

The first objective is to identify the technical, relational, and ethical factors that influence consumer trust during AI-enabled service encounters. Particular attention is given to perceived competence, responsiveness, transparency, fairness, privacy assurance, and access to human support. The study treats these factors as complementary rather than independent because a highly responsive system may still be distrusted if its recommendations appear inaccurate, unfair, opaque, or intrusive.

2.2 To Examine Trust as a Dynamic Service-Encounter Process

The second objective is to conceptualize consumer trust as a judgment that develops and changes across the service journey. Initial trust may depend on organizational reputation, interface presentation, or prior attitudes toward AI, while interaction-based trust develops through service performance. Continued trust is subsequently influenced by consistency, error handling, explanations, data practices, and recovery after service failure.

2.3 To Develop an Integrated Analytical Framework

The third objective is to develop a testable conceptual model connecting major trust-support mechanisms with consumer trust. Because genuine survey or transactional data were not provided, the study uses a transparent simulation to illustrate theoretically expected relationships and establish a methodological foundation for future empirical research.

3. Review of Literature

3.1 Competence, Responsiveness, and Functional Reliability

Consumers are unlikely to develop sustainable trust in an AI-enabled service that repeatedly produces inaccurate information, fails to understand requests, or responds inconsistently. Functional competence therefore forms one of the most fundamental layers of trust.

Glikson and Woolley's review of empirical human-AI trust research emphasizes the importance of reliability, transparency, immediacy, and AI capability in cognitive trust formation.

Research focusing specifically on consumer communication with AI chatbots has likewise identified expertise and responsiveness among the significant positive predictors of trust. Li and colleagues distinguish between chatbot-related factors—including expertise, responsiveness,



anthropomorphism, and ease of use—and broader company- and consumer-related influences such as brand trust, access to human support, perceived risk, and privacy concerns.

Recent qualitative evidence further supports the distinction between functional and relational trust. Wang and colleagues' 2026 customer-service study identified perceived reliability and system competence as a central theme in chatbot trust, with accuracy, responsiveness, transparency, and data security contributing to cognitive confidence. Human-like interaction and emotional connection formed a second, affective dimension.

3.2 Transparency, Fairness, Privacy, and Ethical Assurance

AI-enabled service encounters involve information asymmetry. Consumers usually cannot directly observe training data, model design, decision rules, commercial optimization objectives, or the full extent of data processing occurring behind an interface.

Transparency is frequently proposed as a mechanism for reducing this uncertainty, although contemporary research suggests that appropriate transparency, rather than unlimited technical disclosure, is the more useful concept.

Park and Yoon's experimental research found that algorithm-transparency signals could mitigate negative attitudes toward organizations using AI, particularly in high-involvement situations. Their results also showed that transparency did not uniformly change trust in the AI system itself, demonstrating that organizational trust and system trust should be distinguished.

Research in AI-enabled healthcare applications adds an ethical dimension. A 2025 study reported that transparency influenced trust indirectly through perceived fairness, while privacy concerns altered how fairness was associated with trust. These findings suggest that consumers evaluate not merely whether an AI process is explained but whether its operation appears procedurally legitimate and respectful of personal information.

Broader consumer evidence from the United Kingdom similarly identifies transparency, efficiency, and ethical handling of AI as important influences on consumer trust.

3.3 Anthropomorphism, Service Failure, and Sustained Trust

Human-like characteristics are among the most widely investigated design features of conversational AI. Araujo's experimental research demonstrated that anthropomorphic design cues and communicative framing influence perceptions of conversational agents and the companies deploying them.

Cheng and colleagues further showed that the relationship between chatbot anthropomorphism and consumer trust depends on relationship norms. Consumers do not respond identically to human-like warmth and competence across every interaction context.

The importance of context becomes especially visible when AI systems fail. Errors are unavoidable in complex service environments, making trust recovery as important as initial trust formation.

Gu and colleagues examined sustained consumer trust following AI-chatbot service failures and reported that perceived anthropomorphic characteristics, empathy, and interaction quality could influence how consumers attributed failure. When failures were attributed to inherent AI capability, sustained trust

deteriorated; external attribution could preserve more confidence. AI anxiety further intensified negative trust consequences under internal failure attribution.

Communication style during failure also affects customer reactions. Recent research has therefore moved beyond examining whether a chatbot fails and toward examining how it communicates during and after failure.

4. Research Methodology

4.1 Research Design

The study adopts a conceptual-methodological simulation design.

No human participants were recruited, and the study does not claim to report survey, experimental, transaction, interview, or behavioral data collected from real consumers.

A synthetic analytical population of 600 AI-enabled service-encounter profiles was generated for

4.2 Construct Operationalization

The framework includes six trust antecedents and one principal outcome variable.

Table 1: Operational Framework for AI-Enabled Consumer Trust

Construct	Operational Interpretation	Analytical Scale	Expected Relationship
Perceived Competence	Consumer belief that the AI provides accurate, relevant, and capable assistance	0–100	Positive
Responsiveness	Speed, contextual relevance, and continuity of the AI response	0–100	Positive
Appropriate Transparency	Understandable disclosure concerning AI involvement, reasoning limits, and service process	0–100	Positive/Contextual
Perceived Fairness	Belief that recommendations and decisions are procedurally reasonable and non-arbitrary	0–100	Positive
Privacy Assurance	Confidence that personal data are appropriately handled and protected	0–100	Positive
Human-Support Accessibility	Ability to reach an accountable human when automated service is insufficient	0–100	Positive
Consumer Trust Index	Composite indicator of willingness to rely on the AI-enabled service encounter	0–100	Outcome



The framework deliberately uses the term appropriate transparency rather than absolute transparency because evidence indicates that disclosure effects depend on presentation, timing, context, prior attitudes, and what consumers infer from the disclosure.

4.3 Simulation Procedure

Perceived competence receives the highest illustrative weight because technical reliability represents a foundational condition of cognitive trust, while transparency, fairness, privacy, responsiveness, and human support contribute complementary assurance signals. The model should therefore be interpreted as a reproducible theoretical demonstration rather than a validated predictive instrument.

5. Analysis

5.1 Variation in Consumer Trust Across Service Conditions

The simulation generated a progressive increase in the Consumer Trust Index across the three trust-support conditions. Consumer trust averaged 48.5 in the low-support condition, 67.1 in the moderate condition, and 82.8 in the high condition. The difference between the low- and high-support conditions was therefore 34.3 index points.

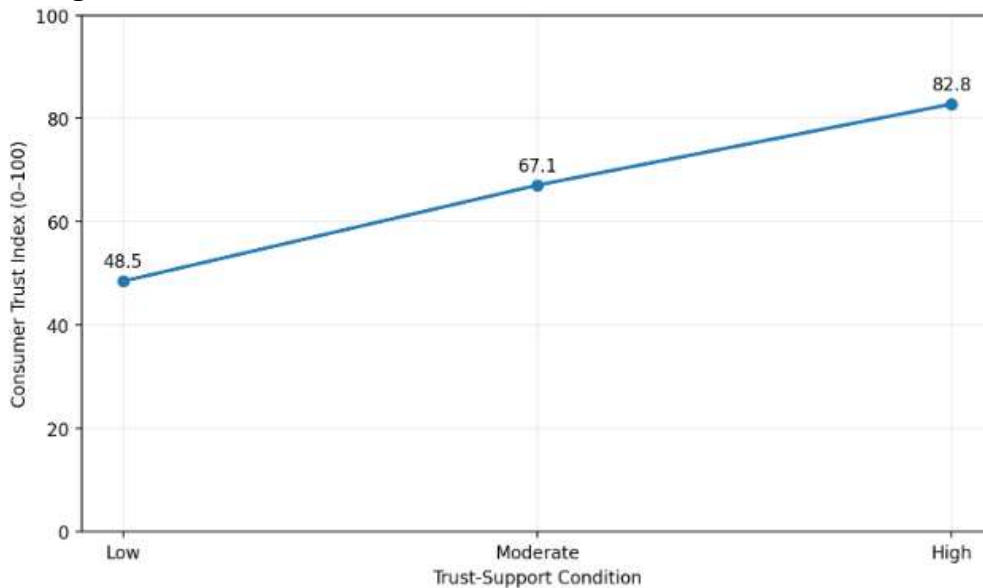
Table 2: Simulated Consumer Trust Outcomes

Analytical Indicator	Low Support	Moderate Support	High Support
Perceived Competence	50.8	69.0	84.1
Responsiveness	52.4	70.4	81.9
Appropriate Transparency	45.6	64.7	80.9
Perceived Fairness	48.8	66.6	82.3
Privacy Assurance	44.6	65.3	82.2
Human-Support Accessibility	47.8	69.1	84.2
Consumer Trust Index	48.5	67.1	82.8

Note: All values are simulation-generated scores. They are not observed consumer responses and must not be interpreted as real-world population estimates.

5.2 Graphical Pattern of Trust Formation

Figure 1: Consumer Trust Formation Across AI Service Conditions



The increase from low to moderate trust support is 18.6 points, while the difference between moderate and high support is a further 15.7 points. The pattern illustrates an important theoretical argument: consumer trust is unlikely to depend on one isolated interface characteristic. Trust strengthens when multiple assurance mechanisms operate together. For example, high competence without privacy assurance may generate useful performance but limited willingness to disclose information. High transparency without competent service delivery may communicate honesty while simultaneously revealing system weakness. Human-like language without reliable problem resolution may generate an initially pleasant encounter but insufficient long-term trust.

6. Discussion

6.1 Trust Should Be Calibrated, Not Maximized

A central contribution of the study is the distinction between high trust and appropriate trust. The managerial objective should not be to persuade consumers to trust an AI system regardless of its actual capability. Excessive consumer trust can become problematic when users over-rely on automated advice, misunderstand model limitations, or assume that human-like conversational fluency guarantees factual accuracy.

Human-AI trust research increasingly distinguishes trust from objective trustworthiness and emphasizes the importance of aligning user confidence with system capability. Organizations should therefore communicate both capability and limitation. A trustworthy AI-enabled service should be able to indicate uncertainty, avoid pretending to possess capabilities it does not have, and escalate appropriate cases rather than continuing an interaction merely to preserve automation rates. Consumer trust should increase because the system behaves dependably—not because consumers are encouraged to underestimate risk.



6.2 The Transparency Paradox in AI Service Encounters

The study also highlights why transparency requires careful implementation.

From an ethical perspective, consumers should know when a materially significant automated system is involved and should receive understandable information about its role. However, empirical evidence demonstrates that disclosure can produce different behavioral responses.

Luo and colleagues found that revealing chatbot identity before an interaction reduced purchases in their field experiment, largely because consumers perceived the disclosed system as less knowledgeable and empathetic. Contemporary evidence offers an additional layer. Park and Yoon found that transparency could signal organizational accountability, particularly under high issue involvement, even where it did not fully overcome distrust of the AI system itself.

6.3 Managerial Implications for AI-Enabled Services

Accurate responses, consistent information, contextual understanding, and reliable execution should be established before investing heavily in avatars, emotional language, or human-like identities. Recent qualitative evidence shows that consumers distinguish emotional-human interaction cues from reliability and competence, while both contribute to trust judgments. Second, fairness and privacy should be incorporated directly into customer-experience design. They should not remain solely technical or compliance responsibilities hidden behind the interface. Research on AI recommendations demonstrates that fairness can function as a pathway connecting transparency and trust, particularly where privacy concerns are salient.

Third, organizations should implement a visible human escalation pathway. Consumers should not be trapped within repetitive automated loops when the AI lacks the information, authorization, emotional sensitivity, or contextual judgment necessary to solve the problem. Fourth, service recovery deserves dedicated design attention. Research on AI-chatbot failure demonstrates that sustained trust depends partly on how consumers interpret the cause of the failure and how effectively the system manages subsequent interaction.

7. Conclusion

Consumer trust has become a fundamental condition for the successful adoption and sustained use of artificial intelligence-enabled service encounters. As AI systems increasingly perform visible customer-facing functions, consumers must evaluate more than technological usefulness. They also need confidence that these systems are competent, fair, transparent, secure, and accountable. Trust therefore represents a multidimensional judgment shaped by system performance, perceived risk, information practices, explainability, and the availability of appropriate safeguards. Building such confidence is essential for organizations seeking to establish reliable and responsible relationships between consumers and AI-enabled services.

The simulation-based analysis illustrates how the proposed framework may respond under different trust-support conditions. The Consumer Trust Index increased from 48.5 under low-support conditions to 67.1 under moderate conditions and 82.8 under high-support conditions. These values are synthetic and are presented solely to demonstrate the behavior of the proposed analytical framework; they should not be interpreted as measured consumer outcomes. The findings nevertheless illustrate the



importance of strengthening transparency, reliability, security, accountability, and user control when designing AI-enabled service environments.

More importantly, the study emphasizes that trust maximization should not be treated as the ultimate objective of AI service design. Organizations should instead pursue calibrated trust, whereby consumer confidence remains proportionate to the system's actual capabilities, limitations, safeguards, and accountability mechanisms. The future of trustworthy AI-enabled services will depend less on making automated systems appear indistinguishable from human employees and more on creating transparent encounters in which consumers understand system behavior, experience dependable performance, retain meaningful control over decisions and personal information, and can access accountable human assistance whenever necessary.

8. References

1. M.-H. Huang and R. T. Rust, "Artificial intelligence in service," *Journal of Service Research*, vol. 21, no. 2, pp. 155–172, 2018, doi: 10.1177/1094670517752459.
2. T. Araujo, "Living up to the chatbot hype: The influence of anthropomorphic design cues and communicative agency framing on conversational agent and company perceptions," *Computers in Human Behavior*, vol. 85, 2018, doi: 10.1016/j.chb.2018.03.051.
3. X. Luo, S. Tong, Z. Fang, and Z. Qu, "Frontiers: Machines vs. humans: The impact of artificial intelligence chatbot disclosure on customer purchases," *Marketing Science*, vol. 38, no. 6, pp. 937–947, 2019, doi: 10.1287/mksc.2019.1192.
4. E. Glikson and A. W. Woolley, "Human trust in artificial intelligence: Review of empirical research," *Academy of Management Annals*, vol. 14, no. 2, pp. 627–660, 2020, doi: 10.5465/annals.2018.0057.
5. X. Cheng, X. Zhang, J. Cohen, and J. Mou, "Human vs. AI: Understanding the impact of anthropomorphism on consumer response to chatbots from the perspective of trust and relationship norms," *Information Processing & Management*, vol. 59, no. 3, Art. no. 102940, 2022, doi: 10.1016/j.ipm.2022.102940.
6. J. Yun *et al.*, "The effects of chatbot service recovery with emotion words on customer satisfaction, repurchase intention and positive word-of-mouth," *Frontiers in Psychology*, vol. 13, 2022, doi: 10.3389/fpsyg.2022.922503.
7. J. Li *et al.*, "Determinants affecting consumer trust in communication with AI chatbots: The moderating effect of privacy concerns," 2023. The study identifies expertise, responsiveness, anthropomorphism, perceived risk, human support, and privacy concerns among relevant trust determinants.
8. d artificial intelligence contexts," *Journal of Retailing and Consumer Services*, vol. 77, Art. no. 103659, 2024, doi: 10.1016/j.jretconser.2023.103659.
9. Y. Li *et al.*, "Developing trustworthy artificial intelligence: Insights from research on interpersonal, human-automation, and human-AI trust," *Frontiers in Psychology*, 2024, doi: 10.3389/fpsyg.2024.1382693.



10. C. Gu *et al.*, “Exploring the mechanism of sustained consumer trust in AI chatbots after service failures: A perspective based on attribution and CASA theories,” *Humanities and Social Sciences Communications*, vol. 11, Art. no. 1400, 2024, doi: 10.1057/s41599-024-03879-5.
11. N. Cai *et al.*, “How the communication style of chatbots influences consumers during failed service experiences,” *Humanities and Social Sciences Communications*, 2024.
12. L. Aquilino *et al.*, “Decoding trust in artificial intelligence: A systematic review,” *Information*, vol. 12, no. 3, Art. no. 70, 2025.
13. B. O’Higgins *et al.*, “Consumer trust in artificial intelligence in the UK and associated drivers,” 2025, doi: 10.1080/23311975.2025.2469765. The study reports important roles for transparency, efficiency, and ethical AI handling in consumer trust.
14. K. Park and H. Y. Yoon, “AI algorithm transparency, pipelines for trust not prisms: Mitigating general negative attitudes and enhancing trust toward AI,” *Humanities and Social Sciences Communications*, vol. 12, Art. no. 1160, 2025.
15. “The impact of AI perceived transparency on trust in AI recommendations in healthcare applications,” *Asia-Pacific Journal of Business Administration*, 2025, doi: 10.1108/APJBA-12-2024-0690.
16. O. Schilke *et al.*, “The transparency dilemma: How AI disclosure erodes trust,” 2025. Experimental evidence indicates that AI disclosure can under some conditions reduce perceived trust and legitimacy.
17. S. Wang, N. Fatima, M. Shahbaz, *et al.*, “Building user trust in AI chatbots for customer service through human-like cues and perceived reliability,” *Scientific Reports*, 2026.
18. M. J. Chae *et al.*, “Too human to trust? How AI human-likeness and context influence consumer trust, satisfaction, and product choices,” 2026.