



International Journal of Business and Management Sciences

2026; 2(4): 172-176 | ISSN(O): 2693-3500

www.sciro.org/ijbms

DOI: [10.5281/zenodo.21991429](https://doi.org/10.5281/zenodo.21991429)

Impact of Role clarity and Trust on Behavioral Intentions among Patients and Healthcare Professionals

Nuzhatul Abrar Siddiqua^{1,*}, Gayathri S², Mohammad Bilal Ahmed³.

¹ Assistant Professor, School of Management, CMR University, Bengaluru.

^{2,3} PG Scholar, Master Business Administration, School of Management, CMR University, Bengaluru.

Email address*: nuzhatul.a@cmr.edu.in

Received: 21.06.2026 | **Revised:** 26.07.2026 | **Accepted:** 07.08.2026

Abstract: Artificial intelligence (AI) adoption in healthcare depends on both service recipients and professionals, yet these groups are rarely examined within one role-sensitive framework. Integrating the Unified Theory of Acceptance and Use of Technology with trust, creepiness, privacy concern, role theory, and complementarity theory, this research reports two field studies from hospitals in Bengaluru, Mumbai, Kolkata, Delhi, and Hyderabad. Study 1 analyzes independent samples exposed to AI in supporting ($n = 508$), augmenting ($n = 513$), and performing ($n = 448$) roles. Partial least squares structural equation modelling and measurement-invariance assessment show that privacy concern and need for human interaction increase creepiness; creepiness reduces trust; and trust, performance expectancy, and social influence increase adoption intention across all roles. Effort expectancy is role contingent. Study 2 analyzes 636 healthcare professionals. Role clarity strongly predicts trust ($\beta = .547$), whereas creepiness reduces trust ($\beta = -.204$) and behavioural intention ($\beta = -.096$). Trust mediates the effects of role clarity and creepiness on intention, and the model explains 59.6% of intention variance. The findings establish trust as the principal coordinating mechanism across the patient–professional AI interface and show that adoption design must reflect what AI is authorized to do, not merely whether AI is present.

Keywords: artificial intelligence, · healthcare services, trust, creepiness, role clarity, UTAUT, PLS-SEM.

1. Introduction

Healthcare AI is moving from isolated prediction tools toward systems that store and interpret records, assist clinical procedures, and perform elements of care delivery. Operational applications can improve monitoring, maintenance, and resource continuity in smart hospitals (Jakata et al., 2025), but technical capability does not ensure use. AI introduces opacity, privacy exposure, uncertain responsibility, and the possibility that patients or professionals will perceive the system as intrusive or unsettling.

Most adoption research applies usefulness- and effort-centred models to a single user group. This leaves two blind spots. First, patients evaluate AI as recipients of a consequential service, while professionals evaluate it as a co-worker embedded in clinical routines. Second, AI can support, augment, or perform service tasks; these roles alter visibility, human control, and perceived responsibility. Treating all AI encounters as equivalent may therefore conceal the conditions under which familiar adoption predictors matter.

We develop a dual-perspective, role-contingent explanation of healthcare AI adoption. Study 1 compares patients and attendants across three independently sampled AI-role scenarios. Study 2 tests how role clarity and creepiness shape professionals' intention directly and through trust. The contribution is a coordinated account of adoption across both sides of the care encounter: trust translates emotional and organizational evaluations into willingness to use AI, while role design determines the strength of that translation.

2. Theoretical framework and hypotheses

2.1. Patient adoption across AI roles

UTAUT explains adoption through performance expectancy, effort expectancy, social influence, and facilitating conditions (Venkatesh et al., 2003). Healthcare AI also requires constructs that capture its adaptive and data-intensive character. Privacy

concern should heighten creepiness because sensitive data collection increases uncertainty about observation and misuse. A need for human interaction should likewise heighten creepiness when an AI interface appears to displace interpersonal reassurance. Creepiness, in turn, should erode trust; trust should increase behavioural intention.

- H1–H3: Privacy concern and need for human interaction are positively related to creepiness, and creepiness is negatively related to trust.
- H4–H7: Trust, performance expectancy, and social influence are positively related to behavioural intention; perceived effort is negatively related to intention.
- H8: Structural relationships differ across supporting, augmenting, and performing AI roles.

2.2. Professional adoption: role clarity and trust

For healthcare professionals, adoption is also an issue of work allocation. Role theory emphasizes clear expectations, and complementarity theory emphasizes a coherent division of contributions between interdependent actors. When responsibility between clinician and AI is explicit, professionals can evaluate competence and accountability more confidently. Role clarity should therefore increase trust and intention. Creepiness should reduce both. Trust is expected to carry part of each effect into behavioural intention.

- H9: Role clarity is positively related to professional behavioural intention.
- H10: Creepiness negatively, related to professional behavioural intention.
- H11: Trust mediates the relationships of role clarity and creepiness with behavioural intention.

3. Method

3.1. Design, setting, and participants

Both studies used cross-sectional field surveys and convenience sampling because a complete sampling frame was unavailable. Data were collected at hospitals in five Indian metropolitan cities: Bengaluru, Mumbai, Kolkata, Delhi, and Hyderabad. Patients in emergency care and their attendants were not approached, identifiable information was not recorded, informed consent was obtained, and institutional ethics approval preceded fieldwork.

The supporting scenario described AI-enabled record storage and health-risk forecasting. The augmenting scenario described surgeon-controlled robotic assistance. The performing scenario described a carebot that monitors vitals, administers medication, and alerts professionals. Study 2 sampled healthcare professionals working with the prospect of AI-enabled care.

3.2. Measures and analysis

Multi-item measures were adapted from established scales. Study 1 measured privacy concern, need for human interaction, creepiness, trust, performance expectancy, perceived effort, social influence, and behavioural intention. Study 2 measured role clarity, creepiness, trust, performance expectancy, effort expectancy, social influence, facilitating conditions, and behavioural intention. SmartPLS 4 was used. Internal consistency, convergent validity, discriminant validity, and collinearity were assessed before bootstrapped structural tests. MICOM preceded multigroup comparisons in Study 1; only partial invariance was obtained, so group-specific estimates are reported.

Table 1: Samples and analytical design

Study	Perspective / condition	Collected	Retained	Analysis
1	AI supporting role (AIS)	638	508	PLS-SEM; MICOM; MGA
1	AI augmenting role (AIA)	680	513	PLS-SEM; MICOM; MGA
1	AI performing role (AIP)	678	448	PLS-SEM; MICOM; MGA
2	Healthcare professionals	680	636	PLS-SEM; mediation

Table 2: Measurement-quality summary

Evidence	Study 1	Study 2	Assessment
Composite reliability	.832–.943	.860–.951	Above .70
Average variance extracted	.646–.846	.608–.830	Above .50
Indicator VIF	< 5	< 5	No severe collinearity
Discriminant validity	Fornell–Larcker, HTMT, cross-loadings	Fornell–Larcker, HTMT, cross-loadings	Established
Measurement invariance	Partial across AI-role groups	Not applicable	Separate group inference

4. Results

4.1. Patient-side structural results

Creepiness consistently undermines trust and is strongest when AI operates in the less visible supporting role (AIS : $\beta = -.626$).

Trust is most strongly associated with intention when AI performs care tasks (AIP : $\beta = .359$).

Perceived effort is significant in supporting and performing conditions but only marginal in the augmenting condition ($\beta = -.070$, $p = .050$). This is substantively plausible: visible clinician control in an augmenting encounter reduces the patient's need to evaluate personal operating effort.

The original report notes that MICOM established partial rather than full invariance. The multigroup evidence should therefore be described cautiously: role-based coefficients differ in magnitude, but most pairwise MGA differences are not statis-

tically significant. The defensible conclusion is role-contingent patterning, not universal moderation of every path.

4.2. Professional-side structural and mediation results

Role clarity is the dominant antecedent of professional trust ($\beta = .547$), and social influence is the largest direct predictor of intention ($\beta = .244$). The model explains 46.7% of trust and 59.6% of behavioural-intention variance. Effort expectancy is the only nonsignificant direct path.

Table 3: Study 1 bootstrapped path estimates by AI role

Relationship	AIP β	AIA β	AIS β	Inference
Privacy concern $\rightarrow$ creepiness	.183***	.272***	.325***	Supported in all roles
Need for human interaction $\rightarrow$ creepiness	.381***	.300***	.375***	Supported in all roles
Creepiness $\rightarrow$ trust	-.470***	-.506***	-.626***	Supported in all roles
Performance expectancy $\rightarrow$ intention	.175**	.144**	.196***	Supported in all roles
Perceived effort $\rightarrow$ intention	-.148**	-.070†	-.204***	Weakest under augmentation
Social influence $\rightarrow$ intention	.204***	.265***	.296***	Supported in all roles
Trust $\rightarrow$ intention	.359***	.336***	.237***	Supported in all roles

Table 4: Study 2 structural model

Path	β	t	p	95% CI / decision
Creepiness $\rightarrow$ intention	-.096	2.121	.034	[-.189, -.010]
Creepiness $\rightarrow$ trust	-.204	5.048	< .001	[-.283, -.123]
Facilitating conditions $\rightarrow$ intention	.146	3.043	.002	[.052, .239]
Performance expectancy $\rightarrow$ intention	.172	3.302	.001	[.074, .276]
Effort expectancy $\rightarrow$ intention	-.078	1.766	.077	[-.164, .010]; not supported
Role clarity $\rightarrow$ intention	.115	2.794	.005	[.035, .197]
Role clarity $\rightarrow$ trust	.547	14.477	< .001	[.470, .619]
Social influence $\rightarrow$ intention	.244	4.948	< .001	[.148, .341]
Trust $\rightarrow$ intention	.108	2.275	.023	[.013, .197]

Table 5: Trust mediation among healthcare professionals

Indirect path	β	t	p	95% CI	Interpretation
Creepiness $\rightarrow$ trust $\rightarrow$ intention	-.022	2.037	.042	[-.044, -.003]	Significant complementary mechanism
Role clarity $\rightarrow$ trust $\rightarrow$ intention	.059	2.240	.025	[.007, .110]	Significant partial mediation

5. Discussion

5.1. Trust as a cross-interface coordinating mechanism

The two studies converge on trust but locate its origins differently. Patients infer trust from the emotional quality of the encounter: privacy concerns and a desire for human contact heighten creepiness, which reduces trust. Professionals infer trust primarily from organizational intelligibility: a clear division of responsibility between clinician and AI produces the strongest path in their model. Trust therefore coordinates adoption across the care interface, but it is built through reassurance for recipients and governance clarity for professionals.

This extends UTAUT by showing that performance and social influence remain relevant, yet are insufficient for adaptive healthcare technologies. Recent evidence similarly places trust,

explanation, and clinical alignment at the centre of medical-AI acceptance (Shevtsova et al., 2024). The present results add role structure: adoption depends not only on what an AI system can do, but on whether users understand its authority, visibility, and relationship with human judgment.

5.2. Role visibility and effort expectancy

The effort result resolves an apparent inconsistency. Patient effort is consequential when AI works behind the scenes or performs care, but weak when a visible clinician operates an augmenting system. Among professionals, effort expectancy is also nonsignificant once role clarity, infrastructure, trust, performance, and social influence are modeled together. In high-stakes settings, clarity of accountability may dominate interface convenience.

5.3. Theoretical contribution

First, the paper integrates technology adoption, role theory, complementarity, and affective responses within one two-sided service framework. Second, independent role samples prevent the AI-role comparison from being reduced to a single generic

vignette. Third, the professional mediation model identifies trust as the mechanism through which work design becomes adoption intention. The result reframes role clarity from an administrative condition into a trust-producing governance signal.

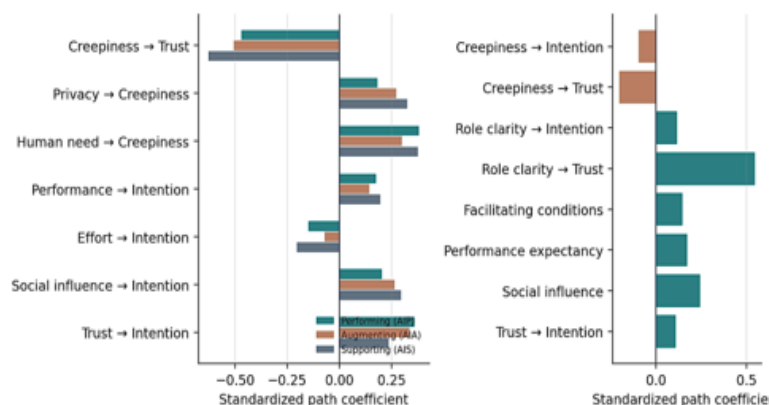


Figure 1: Uniform comparison of Study 1 role-specific paths and Study 2 professional paths. AIA effort expectancy is marginal ($p=.050$); professional effort expectancy is nonsignificant

6. Practical implications

Hospitals should implement role-specific adoption protocols. Supporting AI requires transparency about hidden data processing and escalation. Augmenting AI requires visible confirmation of clinician control. Performing AI requires explicit boundaries, fail-safe handover, and continuous human availability. Patient communication should specify what data are collected, why they are needed, how decisions can be contested, and when a person intervenes.

For professionals, role charters should define task ownership, override authority, accountability, and exception handling before deployment. Training should combine technical operation with scenario-based responsibility exercises. Because facilitating conditions and social influence are significant, peer champions, leadership endorsement, help desks, and protected learning time are likely to matter as much as usability.

Developers should evaluate creepiness, trust, and role clarity during design validation rather than after implementation. Explainability should be matched to the user: patients need understandable consequences and recourse, whereas professionals need clinically actionable reasoning and traceable responsibility.

7. Limitations and future research

The studies are cross-sectional, use convenience samples, and measure behavioural intention rather than observed use. The five-city design improves geographic coverage but does not represent smaller cities or rural health systems. Scenario exposure does not reproduce the emotional or workflow complexity of actual care. Common-method bias cannot be ruled out solely from VIF values, and partial measurement invariance limits strong claims about group differences.

Future work should randomize AI role, explanation, and

human-override cues; link intention to longitudinal usage; and compare public, private, and rural hospitals. Multilevel designs can test whether ethical climate, leadership, professional autonomy, and algorithmic literacy condition trust. Cross-country replication would clarify whether role clarity compensates for cultural differences in uncertainty and authority.

8. Conclusion

Healthcare AI adoption is not a single-user technology problem. It is a coordinated service transformation involving recipients and professionals whose trust is formed through different routes. Patients respond to privacy, human contact, creepiness, and the role AI performs; professionals respond most strongly to role clarity, social influence, infrastructure, and trust. Successful implementation therefore requires role-specific governance and communication, with human accountability made visible throughout the AI-enabled care pathway.

References

1. Ajzen, I. (1985). From intentions to actions: A theory of planned behavior. In J. Kuhl & J. Beckmann (Eds.), *Action control* (pp. 11–39). Springer. https://doi.org/10.1007/978-3-642-69746-3_2
2. Alam, L., & Mueller, S. (2021). Examining the effect of explanation on satisfaction and trust in AI diagnostic systems. *BMC Medical Informatics and Decision Making*, 21, Article 178. <https://doi.org/10.1186/s12911-021-01542-8>
3. Arfi, W. B., Nasr, I. B., Kondrateva, G., & Hikkerova, L. (2021). The role of trust in intention to use the IoT in eHealth. *Technological Forecasting and Social Change*, 165, Article 120479. <https://doi.org/10.1016/j.techfore.2020.120479>
4. Catapan, S. C., Willemann, M. C. A., Calvo, M. C. M., & Bonamigo, E. L. (2024). Trust and artificial intelligence in health care. *Journal of Medical Systems*, 48, Article 21. <https://doi.org/10.1007/s10916-024-02041-0>
5. Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340.

- <https://doi.org/10.2307/249008>
6. Dhagarra, D., Goswami, M., & Kumar, G. (2020). Impact of trust and privacy concerns on technology acceptance in healthcare. *International Journal of Medical Informatics*, 141, Article 104164. <https://doi.org/10.1016/j.ijmedinf.2020.104164>
 7. Gursoy, D., Chi, O. H., Lu, L., & Nunkoo, R. (2019). Consumers acceptance of artificially intelligent device use in service delivery. *International Journal of Information Management*, 49, 157–169. <https://doi.org/10.1016/j.ijinfomgt.2019.03.008>
 8. Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2022). A primer on partial least squares structural equation modeling (PLS-SEM) (3rd ed.). SAGE Publications.
 9. Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science*, 43(1), 115–135. <https://doi.org/10.1007/s11747-014-0403-8>
 10. Hofstede, G. (2011). Dimensionalizing cultures: The Hofstede model in context. *Online Readings in Psychology and Culture*, 2(1). <https://doi.org/10.9707/2307-0919.1014>