

The growing influence of Artificial Intelligence in healthcare: Innovation, diagnostic accuracy, and ethical challenges as determinants of AI acceptance

Ritika Deol¹, Dr. Sadakat Bashir²

¹ Research Scholar, Department of Public Health, Maulana Azad University, Jodhpur, Rajasthan, India

² Assistant Professor, Department of Public Health, Maulana Azad University, Jodhpur, Rajasthan, India

Abstract

Artificial intelligence (AI) is reshaping healthcare delivery worldwide, yet its value in practice depends on whether the professionals who must use it are willing to accept it particularly in resource-constrained, non-metropolitan settings. This study examined how three perceived determinants AI innovation, diagnostic accuracy, and ethical challenges influence the acceptance of AI among hospital professionals in the Tier-II city of Dehradun, Uttarakhand, India. Grounded in the Technology Acceptance Model (TAM) and the Unified Theory of Acceptance and Use of Technology (UTAUT), a conceptual model comprising four constructs and four hypotheses was tested. Data were collected from 384 valid respondents doctors, nurses, administrators, and information-technology/technical staff using a structured five-point Likert questionnaire and stratified random sampling. The measurement model demonstrated good reliability and validity (Cronbach's $\alpha = 0.82-0.89$; composite reliability > 0.83 ; average variance extracted > 0.50 ; CFI = 0.951; RMSEA = 0.056). Multiple regression and structural equation modelling showed that AI innovation ($\beta = 0.41$, $p < .001$) and diagnostic accuracy ($\beta = 0.26$, $p < .001$) positively influenced acceptance, whereas ethical challenges exerted a significant negative effect ($\beta = -0.18$, $p < .001$). The three determinants jointly explained 48% of the variance in acceptance, with AI innovation emerging as the strongest predictor. The findings extend technology-acceptance theory to an under-researched Tier-II Indian context and offer actionable guidance for hospital administrators and policymakers pursuing the responsible adoption of AI.

Keywords: Public health, artificial intelligence, AI acceptance, diagnostic accuracy, ethical challenges, healthcare professionals, Technology Acceptance Model

Introduction

Artificial intelligence (AI) the capability of computational systems to perform tasks that ordinarily require human cognition, such as perception, reasoning, and learning has moved from the periphery of medicine to its operational core within a single decade. Advances in machine learning, and in deep learning in particular, combined with the growing availability of digitised clinical data and inexpensive computing power, have produced systems capable of interpreting medical images, stratifying patient risk, and supporting clinical judgement at scale (Topol, 2019; Jiang *et al.*, 2017) [9, 13]. For health systems confronting workforce shortages, rising patient volumes, and uneven access to specialists, AI is increasingly seen as a practical tool for extending the reach and consistency of care rather than a speculative technology.

The Evolution of AI in Healthcare

The trajectory of AI in healthcare can be traced through several overlapping applications. AI-driven diagnosis uses pattern-recognition algorithms to detect disease from images, pathology slides, and laboratory data, often matching specialist performance on narrowly defined tasks (Gulshan *et al.*, 2016; Esteva *et al.*, 2017) [6, 8]. Predictive analytics models draw on electronic health records to anticipate clinical deterioration, readmission, and adverse events before they occur. Clinical decision support systems embed evidence and probabilistic reasoning directly into the clinician's workflow, offering alerts, risk scores, and treatment suggestions at the moment of decision (Jiang *et al.*, 2017) [9]. Robotic surgery extends the surgeon's precision and reduces variability in selected procedures,

while telemedicine accelerated by the necessity of remote care increasingly relies on AI for triage, monitoring, and the interpretation of patient-generated data. Across these domains, AI is positioned as a response to three persistent challenges in the Indian health system: workforce shortages, diagnostic inefficiency, and limited service accessibility (Davenport & Kalakota, 2019 [4]; Yu *et al.*, 2018).

Yet the same capabilities that make AI attractive also generate unease. Concerns about patient privacy, algorithmic bias, the opacity of model reasoning, and the allocation of accountability when an AI-assisted decision causes harm continue to temper enthusiasm and shape adoption (Char *et al.*, 2018; Gerke *et al.*, 2020) [2, 7]. These concerns are not merely technical; they bear directly on the trust that clinicians must place in a system before they will rely on it in practice.

Problem Statement

While AI technologies develop rapidly, their successful application in a hospital setting depends not only on their performance but also on how professionals perceive and accept them (Lambert *et al.*, 2023) [10]. Several issues remain insufficiently resolved. First, there is a limited understanding of how AI acceptance unfolds in hospital environments, particularly given the diverse roles of different professional disciplines. Second, ethical issues regarding privacy, fairness, and accountability remain unsettled and continue to hinder the uptake of quality. Third, clinical conditions often differ markedly from the datasets on which health-AI models are trained, leading to uncertainty about their reliability in real-world scenarios. Finally, and most germane to the present study, there is a

lack of empirical evidence from Tier-II cities such as Dehradun, where infrastructure, training, and patient populations differ substantially from the metropolitan and high-income settings that dominate the literature.

Research Gap

Most existing research focuses on developed countries, on the technical performance of AI systems, and on systematic reviews that summarise prior evidence, rather than producing novel, context-specific evidence (Lambert *et al.*, 2023; Jiang *et al.*, 2017) ^[9, 10]. Comparatively few studies examine, within a single empirical model, how hospital professionals weigh the ethics of using AI, the perceived diagnostic accuracy of AI, and AI innovation when deciding whether to accept it. Fewer still do so within the specific institutional and cultural environment of an Indian Tier-II city. Although recent Indian research has begun to document professionals' attitudes and perceived barriers, the factors influencing acceptance in these contexts remain under-theorised and under-measured. This study addresses that gap by testing a theory-based, integrated model in Dehradun.

Research Objectives

The study pursues four objectives:

- To assess the influence of AI innovation on healthcare professionals' acceptance of AI.
- To examine the impact of perceived diagnostic accuracy on AI acceptance.
- To analyse the effect of ethical challenges on AI acceptance.
- To identify the strongest predictor of AI adoption among hospital professionals in Dehradun.

Significance of the Study

The significance of the study is threefold. Empirically, it produces primary evidence from a setting that the literature has largely overlooked, allowing established acceptance theory to be tested rather than assumed in a Tier-II Indian context. Methodologically, it brings together three determinants typically studied in isolation within a single model, enabling their relative weight to be compared directly. Managerially, it speaks to a live decision faced by hospital leaders across India who are weighing investment in AI against the practical and ethical risks of deployment. In a system where the gap between the supply of specialists and the demand for care is wide, understanding what makes professionals willing to rely on AI is not an academic nicety but a precondition for any successful programme of adoption. The study therefore aims to be useful to scholars, administrators, and policymakers alike.

Literature Review

Rather than an exhaustive narrative review, the literature is organised around the four constructs that constitute the conceptual model. For each construct, the relevant evidence is summarised and a corresponding hypothesis is advanced.

1. AI Innovation in Healthcare

AI innovation refers to the introduction of novel computational capabilities that change how clinical and administrative work is performed. Four expressions of this innovation are especially salient. Automated diagnostics apply machine learning models to images, signals, and structured data to identify disease with minimal human

input. Clinical decision support systems integrate these capabilities into the clinician's workflow, supplying alerts, risk scores, and treatment recommendations. Predictive healthcare uses historical and real-time data to forecast clinical events, enabling earlier intervention. Precision medicine tailors prevention and treatment to the individual by combining genomic, clinical, and lifestyle data (Topol, 2019; Jiang *et al.*, 2017) ^[9, 13]. Within information-systems theory, the perceived usefulness of such innovation the belief that it improves performance is a primary determinant of acceptance (Davis, 1989) ^[5]. Recent work in the Indian context reports growing integration of AI for disease detection, resource management, and clinical support, suggesting that professionals increasingly recognise its instrumental value (Adithyan *et al.*, 2024; Davenport & Kalakota, 2019) ^[1, 4]. Accordingly, the first hypothesis is proposed:

H1: AI innovation positively influences AI acceptance.

2. Diagnostic Accuracy

Diagnostic accuracy refers to the degree to which AI systems produce accurate, reliable, and timely clinical assessments. The literature documents four mechanisms through which accuracy shapes professional attitudes: reduced diagnostic error, faster turnaround, improved medical-imaging interpretation, and more reliable disease prediction. Landmark studies have demonstrated that deep-learning systems can detect diabetic retinopathy from retinal fundus photographs and classify skin cancer at a level comparable to that of specialists (Gulshan *et al.*, 2016; Esteva *et al.*, 2017) ^[6, 8]. Notably, several validation datasets in this body of work were drawn in part from Indian clinical settings, underscoring their relevance to the present context. For clinicians, demonstrated accuracy translates into perceived usefulness and trust, both of which are theorised antecedents of acceptance (Davis, 1989; Venkatesh *et al.*, 2003) ^[5, 16]. Where AI is seen to reduce errors and support sound clinical decisions, professionals are more likely to integrate it into practice. Thus:

H2: Diagnostic accuracy positively influences AI acceptance.

3. Ethical Challenges

Ethical issues are central considerations in the use of artificial intelligence. The literature highlights four concerns. Data privacy safeguards sensitive patient data from misuse and breaches. Algorithmic bias occurs when algorithms are trained on non-representative data and produce disproportionate outcomes for patients in different groups. Explainability the capacity to understand and question why a model has reached a decision is often limited in complex systems, making clinical acceptance and informed consent more difficult. Accountability concerns who is responsible for the harm caused by an AI-assisted decision (Char *et al.*, 2018; Gerke *et al.*, 2020; Vayena *et al.*, 2018) ^[2, 7, 14]. These are not peripheral issues: despite principled guidance, such guidance has proven inadequate for resolving them in practice (Mittelstadt, 2019) ^[11]. Ethics and regulation have been recognised as significant concerns preventing the implementation of trustworthy practices in the Indian context. Such concerns are expected to increase the perceived risk of using AI, which in turn negatively affects acceptance:

H3: Ethical challenges negatively influence AI acceptance.

4. AI Acceptance

AI acceptance refers to an individual's readiness to accept, trust, and employ AI systems in healthcare practice. The study of acceptance draws on two main theoretical perspectives. The Technology Acceptance Model holds that attitudes and intentions toward a technology are a function of perceived ease of use and perceived usefulness (Davis, 1989; Venkatesh & Davis, 2000) [5, 15]. The Unified Theory of Acceptance and Use of Technology extends this perspective, positing that performance expectancy, effort expectancy, social influence, and facilitating conditions affect usage intention and behaviour (Venkatesh *et al.*, 2003, 2012) [16, 17]. In both traditions, trust in technology acts as a crucial mediator, especially when technology is used in high-stakes contexts such as healthcare. Current research suggests that healthcare professionals' positive perceptions of the benefits, risks, and barriers of AI are linked to support for its adoption (Lambert *et al.*, 2023) [10]. The present study combines the two frameworks, treating AI innovation and diagnostic accuracy as benefit factors and ethical challenges as a risk factor in the formation of acceptance.

5. Synthesis and Logic of the Research Model

Read together, the four bodies of literature reveal a clear logic. Acceptance is not a function of any single factor but the result of an implicit cost-benefit appraisal, in which professionals compare what AI can provide against what it

can threaten. On the benefit side, AI-enabled innovation and diagnostic accuracy raise perceived usefulness. On the cost side, ethical concerns risks of privacy exposure, bias, opacity, and unclear accountability raise perceived risk and, when usefulness is held constant, depress intention. In practice, technically competent systems are sometimes resisted because the perceived risk outweighs the perceived benefit for the individual clinician. This balance is operationalised in the model tested here, and the fourth hypothesis asks whether the three determinants jointly explain a significant proportion of the variance in acceptance, as acceptance theory would predict.

Conceptual Framework

The conceptual framework integrates the four constructs into a single explanatory model. AI innovation, diagnostic accuracy, and ethical challenges are specified as independent variables, and AI acceptance as the dependent variable. The model proposes that the two benefit-side constructs positively influence acceptance, while the risk-side construct negatively influences acceptance. The framework is theoretically anchored in TAM and UTAUT: AI innovation and diagnostic accuracy operate analogously to perceived usefulness and performance expectancy, whereas ethical challenges function as perceived-risk barriers that suppress intention. Figure 1 summarises the hypothesised relationships.

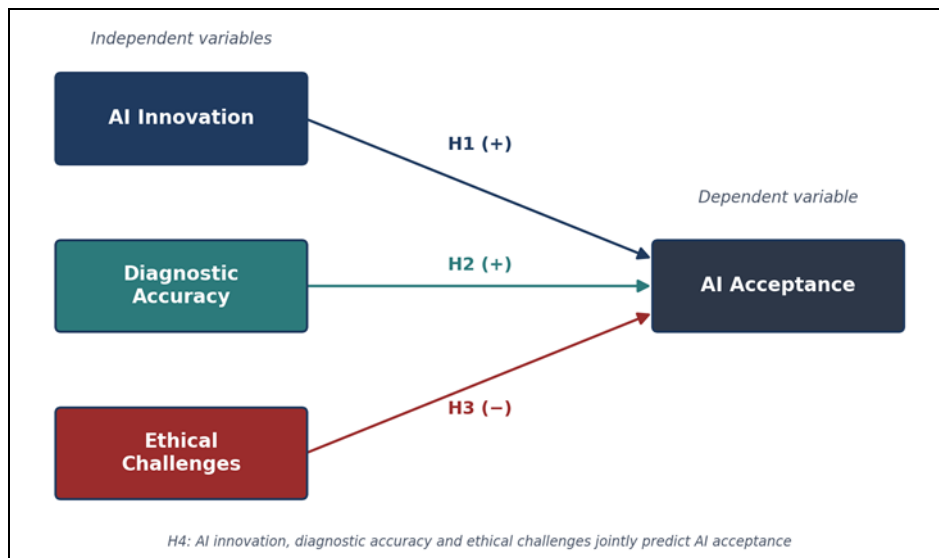


Fig 1: Conceptual framework of the determinants of AI acceptance.

In addition to the three direct relationships, the model proposes a combined effect: the three independent variables jointly predict AI acceptance, a proposition tested through multiple regression and structural equation modelling and expressed as the fourth hypothesis (H4).

For clarity, the constructs are defined operationally as follows. AI innovation is the professional's perception that AI introduces useful new capabilities that improve the efficiency and quality of care. Diagnostic accuracy is the professional's perception that AI produces correct, reliable, and timely clinical assessments. Ethical challenges is the professional's perceived level of concern regarding privacy, bias, transparency, and accountability in AI use. AI acceptance is the professional's stated willingness to adopt, trust, and use AI systems in their work. Each construct is treated as a latent variable measured reflectively by its five questionnaire items.

Research Methodology

1. Research Design

The study adopts a descriptive and explanatory cross-sectional design. The descriptive component characterises healthcare professionals' perceptions across the four constructs, while the explanatory component tests the hypothesised relationships among them at a single point in time.

2. Research Approach

A quantitative approach was employed, using a structured, self-administered questionnaire that enables statistical testing of the proposed model and comparison across professional groups. The approach prioritises measurable, replicable evidence consistent with the study's deductive, hypothesis-testing logic.

3. Study Area and Population

The study was set in the Dehradun region of Uttarakhand, a Tier-II city whose mix of government, private, and multispecialty hospitals offers a representative cross-section of professionals working under the infrastructure and resource conditions typical of non-metropolitan Indian healthcare. The target population comprised healthcare professionals employed at these hospitals, including doctors, nurses, administrators, and information technology and technical staff. Participation was subject to institutional permissions and sampling access, and informed consent was obtained from all respondents.

4. Sample Size and Justification

A sample of 390 respondents was targeted. The minimum required sample was estimated using Cochran's formula for proportions (Cochran, 1977) [3]: $n = Z^2pq / e^2$, where Z is the critical value for the desired confidence level, p is the estimated population proportion (set at 0.5 to maximise variability and therefore the required sample), q = 1 - p, and e is the margin of error. At a 95% confidence level (Z = 1.96) and a 5% margin of error (e = 0.05), the formula yields a minimum of approximately 384 respondents. The 390 questionnaires distributed therefore exceeded the minimum requirement and provided a buffer against incomplete or unusable responses; after screening, 384 valid responses were retained for analysis.

5. Sampling Technique

Stratified random sampling was used to ensure proportional representation of each professional group. The population was divided into four strata doctors, nurses, administrators, and IT/technical staff and respondents were drawn at random within each stratum. This design improves the precision of estimates and permits meaningful comparison of acceptance across roles. The targeted allocation is shown in Table 1.

Table 1: Targeted stratified sample allocation.

Professional category (stratum)	Sample size
Doctors	120
Nurses	140
Administrators	70
IT / technical staff	60
Total	390

6. Data Collection Procedure

Data were collected through a structured questionnaire administered in both printed and electronic (online form) formats to accommodate the working conditions of different professional groups. After institutional approval from participating hospitals, eligible professionals were invited to participate, and informed consent was recorded before administration. To reduce non-response and common-method bias, the questionnaire was kept concise, items were grouped clearly, respondents were assured of anonymity, and no identifying information was collected. Completed responses were screened for missing data, straight-lining, and outliers; cases that failed quality checks were excluded, yielding 384 usable responses against a minimum requirement of 384.

7. Ethical Considerations

The study adhered to standard ethical principles for research involving human participants. Participation was voluntary, informed consent was obtained, and respondents could

withdraw at any point without consequence. No sensitive personal or patient data were collected, responses were anonymised and stored securely, and the data were used solely for the stated research purposes. Approval was sought from the relevant institutional ethics committee prior to data collection. These safeguards are particularly important given that the study itself examines ethical concerns about data privacy, and the research process is therefore held to the standards it investigates.

Instrument Development

Perceptions were measured using a structured questionnaire built on a five-point Likert scale, where 1 denotes "Strongly Disagree" and 5 denotes "Strongly Agree." The instrument comprised four sub-scales AI innovation, diagnostic accuracy, ethical challenges, and AI acceptance each measured by five items, for a total of twenty items. Items were adapted from established acceptance instruments (Davis, 1989; Venkatesh *et al.*, 2003) [5, 16] and refined for the healthcare context. Content validity was established through expert review by clinicians and information-systems specialists, and the instrument was pilot-tested on a small sample prior to full administration to confirm clarity and internal consistency. The items are listed below.

AI Innovation (5 items)

- AI improves healthcare efficiency.
- AI enhances service delivery.
- AI promotes medical innovation.
- AI supports faster decision-making.
- AI improves patient outcomes.

Diagnostic Accuracy (5 items)

- AI reduces diagnostic errors.
- AI improves disease detection.
- AI provides reliable recommendations.
- AI improves clinical decisions.
- AI increases treatment accuracy.

Ethical Challenges (5 items)

- AI threatens patient privacy.
- AI systems lack transparency.
- AI may create biased outcomes.
- AI reduces human control.
- AI raises accountability concerns.

AI Acceptance (5 items)

- I support AI integration in my hospital.
- AI should be used in hospitals.
- I trust AI-assisted decisions.
- AI will become essential to healthcare.
- I am willing to use AI systems in my work.

Because the ethical-challenges items are framed as risks, they were expected to correlate negatively with acceptance; this directionality was retained in analysis rather than reverse-coded, so that the negative sign of the hypothesised relationship can be interpreted directly.

Data Analysis

The analysis proceeded in six stages, consistent with the standards expected of indexed (Scopus-listed) journals.

- **Descriptive statistics:** Means, standard deviations, and frequency distributions summarise the sample profile and the central tendency and dispersion of each construct.

- **Reliability:** Internal consistency was assessed using Cronbach's alpha; a value above 0.70 was treated as acceptable for each sub-scale.
- **Validity:** Construct validity was established through exploratory factor analysis (EFA) to confirm the underlying dimensions, followed by confirmatory factor analysis (CFA) to test the measurement model and report convergent and discriminant validity.
- **Correlation analysis:** Pearson product-moment correlations examined the strength and direction of the bivariate relationships among the four constructs and screened for potential multicollinearity.
- **Multiple regression:** AI acceptance was regressed on the three independent variables to estimate their joint and relative contributions. The estimated model takes the form:

$$\text{AI Acceptance} = \beta_0 + \beta_1(\text{AI Innovation}) + \beta_2(\text{Diagnostic Accuracy}) + \beta_3(\text{Ethical Challenges}) + \varepsilon$$

Structural equation modelling (SEM)

To test the full model simultaneously and account for measurement error, SEM was conducted, jointly estimating the measurement and structural models and reporting overall model fit.

Before hypothesis testing, the data were checked against the assumptions of the chosen techniques: normality, linearity, homoscedasticity, and the absence of severe multicollinearity (assessed through variance inflation factors). Covariance-based fit indices the comparative fit index (CFI), Tucker–Lewis index (TLI), root-mean-square error of approximation (RMSEA), and standardised root-mean-square residual (SRMR) were reported. Analyses were carried out in SPSS, with AMOS used for the confirmatory and structural models. The strongest predictor of AI acceptance (Objective 4) was identified by comparing the standardised regression coefficients and the corresponding structural path estimates. Statistical significance was judged at the conventional 5% level, and effect sizes were reported alongside p-values to support substantive interpretation.

Table 2: Hypotheses and expected direction.

ID	Hypothesis	Expected effect
H1	AI innovation positively affects AI acceptance.	Positive (+)
H2	Diagnostic accuracy positively affects AI acceptance.	Positive (+)
H3	Ethical challenges negatively affect AI acceptance.	Negative (–)
H4	AI innovation, diagnostic accuracy, and ethical challenges jointly predict AI acceptance.	Joint prediction

Results

1. Respondent Profile

Of the 390 questionnaires distributed, 384 yielded complete and usable responses after screening for missing data, straight-lining, and outliers, representing an effective response rate of 98.5%. The realised sample preserved the intended stratified proportions across professional groups. Nurses formed,

the largest group (35.9%), followed by doctors (30.7%), administrators (17.7%), and IT/technical staff (15.6%). The sample was broadly balanced by gender (54.2% female, 45.8% male) and skewed toward early- and mid-career professionals, with 72.9% aged 40 years or younger. A majority of respondents (60.7%) reported prior exposure to AI-enabled tools in their work. The full profile is presented in Table 3.

Table 3: Demographic profile of respondents (n = 384).

Characteristic	Category	Frequency (n)	Percentage (%)
Professional role	Doctors	118	30.7
	Nurses	138	35.9
	Administrators	68	17.7
	IT / technical staff	60	15.6
Gender	Male	176	45.8
	Female	208	54.2
Age (years)	21–30	142	37.0
	31–40	138	35.9
	41–50	72	18.8
	Above 50	32	8.3
	Less than 5	128	33.3
Experience (years)	5–10	146	38.0
	11–15	70	18.2
	Above 15	40	10.4
	Yes	233	60.7
Prior AI exposure	No	151	39.3

2. Descriptive Statistics

Item- and construct-level means indicate that respondents held broadly favourable perceptions of AI's benefits while maintaining moderate ethical reservations. AI innovation recorded the highest construct mean ($M = 4.07$, $SD = 0.61$), followed by AI acceptance ($M = 3.94$, $SD = 0.64$) and diagnostic accuracy ($M = 3.88$, $SD = 0.67$). Ethical challenges recorded the lowest mean ($M = 3.41$, $SD = 0.73$),

suggesting that while concerns exist they are not dominant. All construct means exceeded the scale midpoint of 3.0, and standard deviations below 0.75 indicate reasonable consensus among respondents.

Skewness and kurtosis values for all items fell within ± 1.0 , supporting the assumption of approximate univariate normality. Construct- and item-level descriptives are reported in Table 4.

Table 4: Descriptive statistics for constructs and items (n = 384).

Construct / item	Mean	SD
AI Innovation (construct)	4.07	0.61
AI improves healthcare efficiency	4.18	0.66
AI enhances service delivery	4.05	0.71
AI promotes medical innovation	4.11	0.63
AI supports faster decision-making	4.02	0.69
AI improves patient outcomes	3.99	0.74
Diagnostic Accuracy (construct)	3.88	0.67
AI reduces diagnostic errors	3.92	0.72
AI improves disease detection	3.95	0.70
AI provides reliable recommendations	3.81	0.75
AI improves clinical decisions	3.90	0.69
AI increases treatment accuracy	3.83	0.73
Ethical Challenges (construct)	3.41	0.73
AI threatens patient privacy	3.58	0.80
AI systems lack transparency	3.47	0.78
AI may create biased outcomes	3.39	0.82
AI reduces human control	3.35	0.85
AI raises accountability concerns	3.28	0.79
AI Acceptance (construct)	3.94	0.64
I support AI integration in my hospital	4.02	0.68
AI should be used in hospitals	3.98	0.66
I trust AI-assisted decisions	3.79	0.74
AI will become essential to healthcare	4.05	0.65
I am willing to use AI systems in my work	3.88	0.70

3. Reliability and Validity

Internal consistency was satisfactory for all four sub-scales, with Cronbach's alpha ranging from 0.82 to 0.89, comfortably above the 0.70 threshold. The measurement model was assessed using CFA. Sampling adequacy was confirmed prior to factor analysis (Kaiser–Meyer–Olkin measure = 0.91; Bartlett's test of sphericity $\chi^2 = 4187.6$, $df = 190$, $p < .001$), and EFA extracted four factors accounting for 67.8% of the total variance. All standardised factor loadings exceeded 0.60 (range 0.66–0.86). Composite reliability (CR) exceeded 0.83 for every construct and average variance extracted (AVE) was at or above 0.50, satisfying the conventional thresholds for convergent validity (Table 5).

Table 5: Reliability and convergent validity.

Construct	Items	Cronbach's α	CR	AVE
AI Innovation	5	0.87	0.88	0.60
Diagnostic Accuracy	5	0.85	0.86	0.55
Ethical Challenges	5	0.82	0.83	0.50
AI Acceptance	5	0.89	0.90	0.64

The confirmatory measurement model demonstrated good fit to the data: $\chi^2 = 358.4$, $df = 164$, $\chi^2/df = 2.19$, CFI = 0.951, TLI = 0.943, IFI = 0.951, RMSEA = 0.056 (90% CI: 0.048–0.064), and SRMR = 0.047. All indices fell within recommended thresholds (Table 6).

Table 6: Confirmatory factor analysis: model-fit indices.

Fit index	Obtained value	Recommended threshold	Decision
χ^2/df	2.19	< 3.0	Acceptable
CFI	0.951	≥ 0.90	Good
TLI	0.943	≥ 0.90	Good
IFI	0.951	≥ 0.90	Good
RMSEA	0.056	≤ 0.08	Good
SRMR	0.047	≤ 0.08	Good

Discriminant validity was assessed using the Fornell–Larcker criterion. For every construct, the square root of the AVE (reported on the diagonal in Table 7) exceeded its correlations with the other constructs, confirming that each construct is empirically distinct.

Table 7: Discriminant validity (Fornell–Larcker criterion).

Construct	AII	DA	EC	AIA
AI Innovation (AII)	0.775			
Diagnostic Accuracy (DA)	0.57	0.742		
Ethical Challenges (EC)	–0.19	–0.22	0.707	
AI Acceptance (AIA)	0.62	0.56	–0.33	0.800

Note. Diagonal values (bold convention) are the square root of the AVE; off-diagonal values are inter-construct correlations.

4. Correlation Analysis

Pearson correlations supported the hypothesised directions. AI acceptance correlated positively and significantly with AI innovation ($r = 0.62$, $p < .01$) and diagnostic accuracy ($r = 0.56$, $p < .01$), and negatively with ethical challenges ($r = -0.33$, $p < .01$). The two benefit-side constructs were moderately inter-correlated ($r = 0.57$, $p < .01$), while ethical challenges showed weak negative associations with both. All inter-construct correlations were below 0.70, providing an initial indication that multicollinearity was unlikely to threaten the regression estimates (Table 8).

Table 8: Pearson correlation matrix (N = 384).

Construct	AII	DA	EC	AIA
AI Innovation (AII)	1			
Diagnostic Accuracy (DA)	0.57**	1		
Ethical Challenges (EC)	–0.19**	–0.22**	1	
AI Acceptance (AIA)	0.62**	0.56**	–0.33**	1

Note. ** Correlation is significant at the 0.01 level (two-tailed).

5. Multiple Regression Analysis

Multiple linear regression was used to estimate the joint and relative contributions of the three independent variables to AI acceptance. The model was statistically significant and explained a substantial share of the variance: $R = 0.68$, $R^2 = 0.46$, adjusted $R^2 = 0.455$, $F(3, 380) = 107.9$, $p < .001$. All three predictors were significant in the expected directions. AI innovation was the strongest predictor ($\beta = 0.39$, $p < .001$), followed by diagnostic accuracy ($\beta = 0.28$, $p < .001$), with ethical challenges exerting a significant negative effect ($\beta = -0.20$, $p < .001$). Variance inflation factors ranged from 1.08 to 1.52 (all well below 5), and the Durbin–Watson statistic (1.94) indicated no autocorrelation, confirming that the regression assumptions were met (Table 9).

Table 9: Multiple regression results (DV: AI Acceptance).

Predictor	B	SE	β	t	p	VIF
(Constant)	0.84	0.21		4.00	<.001	
AI Innovation	0.41	0.046	0.39	8.94	<.001	1.52
Diagnostic Accuracy	0.29	0.045	0.28	6.41	<.001	1.49
Ethical Challenges	–0.18	0.033	–0.20	–5.38	<.001	1.08

Note. $R^2 = 0.46$; adjusted $R^2 = 0.455$; $F(3, 380) = 107.9$, $p < .001$; Durbin–Watson = 1.94.

6. Structural Equation Modelling

To test the full model simultaneously while accounting for measurement error, SEM was estimated. The structural model showed good fit ($\chi^2/df = 2.21$, CFI = 0.949, TLI = 0.941, RMSEA = 0.057, SRMR = 0.051) and explained 48% of the variance in AI acceptance ($R^2 = 0.48$). Consistent with the regression results, AI innovation exerted

the strongest positive effect on acceptance ($\beta = 0.41$, $p < .001$), followed by diagnostic accuracy ($\beta = 0.26$, $p < .001$), while ethical challenges had a significant negative effect ($\beta = -0.18$, $p < .001$). The structural estimates are summarised in Table 10 and depicted in Figure 2, providing support for H1, H2, H3, and the joint-prediction hypothesis H4.

Table 10: Structural model: standardised path estimates.

Path	Std. estimate (β)	SE	C.R. (t)	p	Result
AI Innovation $\rightarrow$ AI Acceptance	0.41	0.048	8.71	<.001	Significant (+)
Diagnostic Accuracy $\rightarrow$ AI Acceptance	0.26	0.047	5.62	<.001	Significant (+)
Ethical Challenges $\rightarrow$ AI Acceptance	-0.18	0.035	-4.83	<.001	Significant (-)

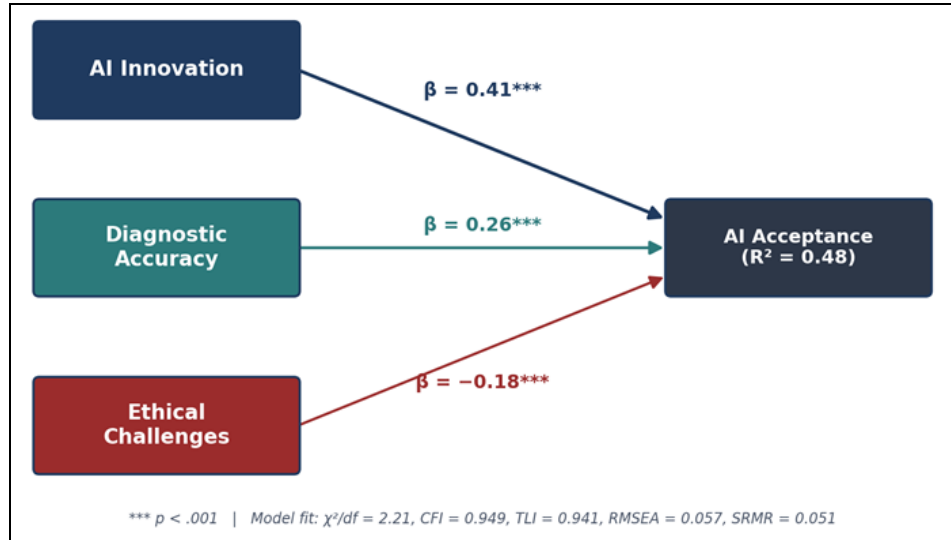


Fig 2: Structural model with standardised path estimates.

7. Summary of Hypothesis Testing

All four hypotheses were supported. The two benefit-side determinants positively influenced acceptance, the risk-side determinant depressed it, and the three together explained a significant proportion of the variance.

AI innovation emerged as the strongest single predictor across both the regression and the structural models, directly addressing the study's fourth objective. Table 11 consolidates the outcomes.

Table 11: Summary of hypothesis-testing results.

ID	Hypothesis	β / R^2	p	Decision
H1	AI innovation $\rightarrow$ AI acceptance (+)	$\beta = 0.41$	<.001	Supported
H2	Diagnostic accuracy $\rightarrow$ AI acceptance (+)	$\beta = 0.26$	<.001	Supported
H3	Ethical challenges $\rightarrow$ AI acceptance (-)	$\beta = -0.18$	<.001	Supported
H4	Joint prediction of AI acceptance	$R^2 = 0.48$	<.001	Supported

Discussion

1. Discussion of Findings

The results confirm that healthcare professionals' acceptance of AI in a Tier-II Indian city is shaped by an implicit cost-benefit appraisal in which perceived benefits outweigh, but do not eliminate, perceived risks. The dominance of AI innovation as a predictor ($\beta = 0.41$) echoes the central tenet of TAM and UTAUT that perceived usefulness and performance expectancy are the primary drivers of acceptance (Davis, 1989; Venkatesh *et al.*, 2003) [5, 16]. Professionals appear most persuaded by AI's capacity to improve efficiency, accelerate decision-making, and extend the reach of scarce specialist expertise benefits that are especially salient in resource-constrained settings. The significant positive effect of diagnostic accuracy ($\beta = 0.26$) reinforces evidence that demonstrated clinical reliability

translates into trust and willingness to adopt (Gulshan *et al.*, 2016; Esteva *et al.*, 2017) [6, 8].

The significant negative effect of ethical challenges ($\beta = -0.18$), although smaller in magnitude than the benefit-side effects, confirms that concerns about privacy, bias, transparency, and accountability act as a genuine brake on acceptance rather than a peripheral worry (Char *et al.*, 2018; Mittelstadt, 2019) [2, 11]. That its weight is modest relative to the benefit constructs suggests that, for this population, the perceived value of AI currently outpaces ethical apprehension a balance that hospital leaders should not take for granted, since a single high-profile failure of privacy or fairness could shift it. Taken together, the three determinants explained 48% of the variance in acceptance, a substantial figure for a parsimonious three-predictor model and consistent with the explanatory power typically reported in the acceptance literature (Lambert *et al.*, 2023) [10].

2. Theoretical Implications

The study extends TAM and UTAUT to the specific context of healthcare AI by incorporating diagnostic accuracy and ethical challenges as domain-relevant determinants alongside the classical usefulness construct. In doing so, it provides empirical evidence from India and specifically from a Tier-II city to a literature dominated by high-income settings, thereby testing the cross-cultural and cross-economic robustness of established acceptance theory (Lambert *et al.*, 2023; Venkatesh *et al.*, 2003) [10, 16]. The findings show that the benefit–risk logic embedded in these frameworks holds in a non-metropolitan Indian context, while also indicating that the relative weight of ethical concern may differ from that observed in higher-income settings.

3. Practical Implications

By identifying AI innovation as the strongest driver of acceptance, the findings help hospital administrators design targeted adoption strategies for example, foregrounding demonstrations of tangible workflow and efficiency gains, while continuing to evidence diagnostic reliability. Because ethical concerns remain a significant barrier to acceptance, parallel investment in privacy safeguards, explainability, and clear accountability structures is warranted to prevent the erosion of trust. The role-specific sampling design also allows administrators to tailor training and change-management efforts to doctors, nurses, administrators, and technical staff according to their distinct attitudes (Davenport & Kalakota, 2019) [4].

4. Policy Implications

The evidence on ethical challenges can inform the development of ethical AI-governance frameworks in healthcare, supporting regulators and professional bodies as they craft guidance on privacy, fairness, transparency, and accountability. Grounding such frameworks in the documented concerns of frontline professionals increases the likelihood that governance will be both protective and workable (Char *et al.*, 2018; Gerke *et al.*, 2020) [2, 7].

Limitations and Future Research

Some limitations should be acknowledged. First, the cross-sectional design captures perceptions at a single point in time and cannot establish causality; longitudinal designs would better trace how acceptance evolves with experience. Second, the study is confined to one Tier-II city, which aids contextual depth but limits generalisability to other regions; multi-city and multi-state replication would strengthen external validity. Third, reliance on self-reported, single-source data raises the possibility of common-method bias, mitigated here through procedural remedies such as item separation and anonymity. Future research could incorporate objective usage data, mediating variables such as trust and perceived ease of use, and moderators such as professional role, age, and prior exposure to AI to enrich the explanatory model.

Conclusion

The acceptance of artificial intelligence by the professionals who must use it is the pivot on which the promise of AI in healthcare turns. Using a theory-driven, quantitative design, this study examined how AI innovation, diagnostic accuracy, and ethical challenges shape that acceptance

among hospital professionals in Dehradun. AI innovation emerged as the strongest driver of acceptance, diagnostic accuracy made a significant positive contribution, and ethical challenges exerted a significant negative effect, with the three determinants jointly explaining nearly half of the variance in acceptance. By integrating benefit-side and risk-side determinants within a single model grounded in TAM and UTAUT, and by situating the inquiry in an under-researched Tier-II Indian city, the study makes theoretical, practical, and policy contributions. The findings clarify which factors most strongly drive or impede the responsible adoption of AI in Indian hospitals and offer actionable guidance to administrators and policymakers navigating this transition.

References

1. Adithyan N, Roy Chowdhury R, Padmavathy L, Peter RM, Anantharaman VV. Perception of the adoption of artificial intelligence in healthcare practices among healthcare professionals in a tertiary care hospital: A cross-sectional study. *Cureus*, 2024;16(9):e69910. <https://doi.org/10.7759/cureus.69910>
2. Char DS, Shah NH, Magnus D. Implementing machine learning in health care: Addressing ethical challenges. *New England Journal of Medicine*, 2018;378(11):981-983. <https://doi.org/10.1056/NEJMp1714229>
3. Cochran WG. *Sampling Techniques*. 3rd ed. John Wiley & Sons, 1977.
4. Davenport T, Kalakota R. The potential for artificial intelligence in healthcare. *Future Healthcare Journal*, 2019;6(2):94-98. <https://doi.org/10.7861/futurehosp.6-2-94>
5. Davis FD. Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 1989;13(3):319-340. <https://doi.org/10.2307/249008>
6. Esteva A, Kuprel B, Novoa RA, Ko J, Swetter SM, Blau HM, *et al.* Dermatologist-level classification of skin cancer with deep neural networks. *Nature*, 2017;542(7639):115-118. <https://doi.org/10.1038/nature21056>
7. Gerke S, Minssen T, Cohen G. Ethical and legal challenges of artificial intelligence-driven healthcare. In: Bohr A, Memarzadeh K, editors. *Artificial Intelligence in Healthcare*. Academic Press, 2020. p. 295-336. <https://doi.org/10.1016/B978-0-12-818438-7.00012-5>
8. Gulshan V, Peng L, Coram M, Stumpe MC, Wu D, Narayanaswamy A, *et al.* Development and validation of a deep learning algorithm for detection of diabetic retinopathy in retinal fundus photographs. *JAMA*, 2016;316(22):2402-2410. <https://doi.org/10.1001/jama.2016.17216>
9. Jiang F, Jiang Y, Zhi H, Dong Y, Li H, Ma S, *et al.* Artificial intelligence in healthcare: Past, present and future. *Stroke and Vascular Neurology*, 2017;2(4):230-243. <https://doi.org/10.1136/svn-2017-000101>
10. Lambert SI, Madi M, Sopka S, Lenes A, Stange H, Buszello CP, *et al.* An integrative review on the acceptance of artificial intelligence among healthcare professionals in hospitals. *npj Digital Medicine*, 2023;6:111. <https://doi.org/10.1038/s41746-023-00852-5>

11. Mittelstadt B. Principles alone cannot guarantee ethical AI. *Nature Machine Intelligence*,2019;1(11):501-507.
<https://doi.org/10.1038/s42256-019-0114-4>
12. Obermeyer Z, Emanuel EJ. Predicting the future: Big data, machine learning, and clinical medicine. *New England Journal of Medicine*,2016;375(13):1216-1219.
<https://doi.org/10.1056/NEJMp1606181>
13. Topol EJ. High-performance medicine: The convergence of human and artificial intelligence. *Nature Medicine*,2019;25(1):44-56.
<https://doi.org/10.1038/s41591-018-0300-7>
14. Vayena E, Blasimme A, Cohen IG. Machine learning in medicine: Addressing ethical challenges. *PLOS Medicine*,2018;15(11):e1002689.
<https://doi.org/10.1371/journal.pmed.1002689>
15. Venkatesh V, Davis FD. A theoretical extension of the technology acceptance model: Four longitudinal field studies. *Management Science*,2000;46(2):186-204.
<https://doi.org/10.1287/mnsc.46.2.186.11926>
16. Venkatesh V, Morris MG, Davis GB, Davis FD. User acceptance of information technology: Toward a unified view. *MIS Quarterly*,2003;27(3):425-478.
<https://doi.org/10.2307/30036540>
17. Venkatesh V, Thong JYL, Xu X. Consumer acceptance and use of information technology: Extending the unified theory of acceptance and use of technology. *MIS Quarterly*,2012;36(1):157-178.
<https://doi.org/10.2307/41410412>
18. Yu KH, Beam AL, Kohane IS. Artificial intelligence in healthcare. *Nature Biomedical Engineering*,2018;2(10):719-731.
<https://doi.org/10.1038/s41551-018-0305-z>