

A moderated mediation model of AI acceptance in healthcare: Perceived AI usage competence, AI trust, and AI anxiety

Sağlık hizmetlerinde yapay zekâ kabulüne ilişkin düzenleyici aracılık modeli: Algılanan yapay zekâ kullanım yetkinliği, yapay zekâya güven ve yapay zekâ kaygısı

Mehmet Kılınç¹ 

¹Assist. Prof. Dr., Bayburt University, Faculty of Health Sciences, Department of Health Management, Bayburt, Türkiye



Abstract

This study examined whether perceived AI usage competence is associated with AI acceptance in healthcare and whether this association is mediated by AI trust and moderated by AI anxiety. A cross-sectional, correlational survey design was used. Data were collected online from 422 adults in Türkiye who had used healthcare services within the previous six months. Perceived AI usage competence was modeled as the independent variable, AI trust as the mediating variable, AI acceptance in healthcare as the outcome variable, and AI anxiety as the moderating variable. Descriptive statistics, reliability analysis, confirmatory factor analysis, Pearson correlation analysis, and Hayes PROCESS Macro Model 8 were used to analyze the data. The findings indicated that perceived AI usage competence was positively associated with AI trust and AI acceptance in healthcare, and AI trust was positively associated with AI acceptance. The trust-mediated indirect association between perceived AI usage competence and AI acceptance was statistically significant. AI anxiety moderated both the competence-trust association and the direct competence-acceptance association, with weaker associations observed at higher levels of AI anxiety. Overall, the findings suggest that AI acceptance in healthcare is related not only to perceived technical competence but also to psychosocial factors such as trust and anxiety.

Keywords: AI acceptance, AI anxiety, AI trust, perceived AI usage competence, healthcare

Öz

Bu çalışma, algılanan yapay zekâ kullanım yetkinliğinin sağlık hizmetlerinde yapay zekâ kabulüyle ilişkili olup olmadığını ve bu ilişkinin yapay zekâ güveni tarafından aracılık edilip edilmediğini ve yapay zekâ kaygısı tarafından düzenlenip düzenlenmediğini incelemektedir. Çalışmada kesitsel, korelasyonel bir anket tasarımı kullanılmıştır. Veriler, son altı ay içinde sağlık hizmetlerinden yararlanmış Türkiye'deki 422 yetişkinden çevrimiçi olarak toplanmıştır. Algılanan yapay zekâ kullanım yetkinliği bağımsız değişken, yapay zekâ güveni aracı değişken, sağlık hizmetlerinde yapay zekâ kabulü sonuç değişkeni ve yapay zekâ kaygısı düzenleyici değişken olarak modellenmiştir. Verilerin analizinde tanımlayıcı istatistikler, güvenilirlik analizi, doğrulayıcı faktör analizi, Pearson korelasyon analizi ve Hayes PROCESS Makro Model 8 kullanılmıştır. Bulgular, algılanan yapay zekâ kullanım yetkinliğinin yapay zekâya yönelik güven ve sağlık hizmetlerinde yapay zekâ kabulü ile pozitif olarak ilişkili olduğunu ve yapay zekâya yönelik güvenin yapay zekâ kabulü ile pozitif olarak ilişkili olduğunu göstermektedir. Algılanan yapay zekâ kullanım yetkinliği ile yapay zekâ kabulü arasındaki güven aracılı dolaylı ilişki istatistiksel olarak anlamlıdır. Yapay zekâ kaygısı, hem yeterlilik-güven ilişkisini hem de doğrudan yeterlilik-kabul ilişkisini düzenlemiş olup, yapay zekâ kaygısının daha yüksek seviyelerinde daha zayıf ilişkiler gözlemlenmiştir. Genel olarak, bulgular yapay zekânın sağlık hizmetlerinde kabulünün yalnızca algılanan teknik yeterlilikle değil, aynı zamanda güven ve kaygı gibi psikososyal faktörlerle de ilişkili olduğunu göstermektedir.

Anahtar Kelimeler: Yapay zekaya güven, YZ kabulü, YZ kaygısı, algılanan YZ kullanım yetkinliği, sağlık hizmetleri

Citation:

Kılınç, M. (2026). A moderated mediation model of AI acceptance in healthcare: Perceived AI usage competence, AI trust, and AI anxiety. *OPUS- Journal of Society Research*, 23, e1922808. <https://doi.org/10.26466/opusjsr.1922808>

Open Access Statement:

This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). This license permits unrestricted use, distribution, and reproduction in any medium, provided that the original author(s) and the source are properly credited. The original publication in *OPUS Journal of Society Research* must be cited in accordance with accepted academic practice.

Review Note:

Evaluated by Double-Blind Peer Review

Ethics Reporting:

To report potential ethical concerns, contact: editorialoffice@opusjournal.net

Similarity screening was conducted via intihal.net.



Introduction

Artificial intelligence (AI) is rapidly becoming widespread in numerous areas of healthcare, including diagnosis, triage, decision support, data management, patient monitoring, and process optimization. However, the high clinical and managerial potential of AI in healthcare alone does not guarantee successful implementation. Current reviews and consensus documents emphasize that, for AI to be used effectively in real-world healthcare applications, reliability, governance, explainability, ethical compliance, and user acceptance must be considered alongside technical performance (Hassan et al., 2024; Lekadir et al., 2025).

The main factors hindering the adoption of AI applications in healthcare include data privacy concerns, legal and ethical uncertainties, concerns about reduced human oversight, the possibility of algorithmic errors, and users' limited knowledge and experience with the technology. Because healthcare decisions directly affect patient safety, the acceptance process for AI is more sensitive than in other sectors. Indeed, current systematic reviews show that the widespread adoption of AI in healthcare is closely related not only to technological infrastructure but also to the perceptions, reservations, and contextual assessments of clinicians and other stakeholders (Sahoo et al., 2025; Scipion et al., 2025).

In this context, trust is one of the most critical determinants of AI acceptance in healthcare. Healthcare users evaluate AI systems not only for usefulness and speed, but also for accuracy, accountability, transparency, privacy, and human oversight. Therefore, trust acts as a bridge between the usability and adoptability of technology. Studies show that trust and acceptance of AI applications are influenced not only by technological characteristics but also by perceptions related to people, institutions, and ethical/legal frameworks. In particular, the TrAAIT model developed by Stevens and Stetson revealed that clinician trust explains a significant portion of the variance in AI adoption (Shevtsova et al., 2024; Stevens & Stetson, 2023).

One key variable in building trust in AI is an individual's competence with AI. User competence refers to an individual's capacity to understand AI tools, identify appropriate use cases, interpret outputs, and interact consciously with technology. The more familiar and proficient a user is with technology, the more predictable, manageable, and useful they perceive it to be. Accordingly, current studies show that as AI literacy and competence with AI increase, attitudes towards AI become more positive. Research conducted on healthcare professionals has also reported strong positive correlations between AI literacy and positive attitudes (Çapuk et al., 2026; Liu et al., 2025).

However, positive cognitive evaluations of AI do not always automatically translate into acceptance. Another concept that has gained prominence in recent years is AI anxiety. This anxiety can include fear of making mistakes, the belief that the human role will diminish, feelings of professional or personal inadequacy, concerns about data security, and perceptions that the technology is uncontrollable. Current findings show that AI training can reduce anxiety; conversely, higher AI anxiety is associated with more negative attitudes towards the technology. Similarly, the literature indicates that anxiety is not only a directly negative emotion but can also be a boundary condition that weakens the positive effects of AI literacy or competence (LaGore et al., 2025; Li & Liu, 2025; Liu et al., 2025).

A review of the existing literature reveals that a significant portion of AI research in healthcare focuses on clinicians, nurses, and other healthcare professionals. However, the acceptance of AI applications among adult healthcare users should also be considered, as healthcare users' perceptions of trust, anxiety, competence, privacy, safety, and human oversight may differ from those of professionals. Therefore, evidence from professional samples should be interpreted cautiously when applied to general healthcare users. Furthermore, existing studies mostly address AI literacy, trust, anxiety, and acceptance variables separately or within the framework of binary relationships; studies that examine these variables together within the same model, in terms of mediation and moderation, are

limited (Scipion et al., 2025; Shevtsova et al., 2024). Accordingly, the findings of the present study should be interpreted in relation to adult individuals in Türkiye who participated in the online survey and used healthcare services within the previous six months, rather than as evidence of societal acceptance in the broader population.

Although a substantial body of research on healthcare AI acceptance has focused on clinicians and healthcare professionals, patient and public acceptance represents a distinct and equally important dimension of implementation. Healthcare users may evaluate AI systems based on different concerns, including perceived diagnostic safety, privacy, human oversight, transparency, and the risk of dehumanized care. Therefore, evidence from professional samples should be interpreted cautiously when applied to general healthcare users. To address this issue, the present study focuses specifically on adult healthcare users and positions AI trust, AI anxiety, and perceived AI usage competence as psychosocial factors relevant to public-facing AI acceptance in healthcare.

Recent patient- and public-oriented studies further indicate that trust and acceptance of healthcare AI are shaped by contextual and governance-related factors (Bracic et al., 2026; Kauttonen et al., 2025; Shevtsova et al., 2024; Švestková et al., 2025). Bracic et al. (2026) showed that patient trust in and preference for medical AI scenarios were associated with AI performance, clinician presence, disclosure of representative data, and systemic governance mechanisms. Similarly, Kauttonen et al. (2025) reported that people's trust in and acceptance of healthcare AI vary across use cases and are influenced by perceived risk and application context. Shevtsova et al. (2024) also emphasized that trust in and acceptance of medical AI are shaped by human-related, technology-related, ethical, and legal factors, including transparency, accountability, data use, and institutional safeguards. In addition, Švestková et al. (2025) found that trust in generative AI for health-related purposes was strongly associated with willingness to use it, highlighting the central

role of trust in public-facing AI adoption. Taken together, these studies suggest that AI acceptance in healthcare should be examined not only through technical competence, but also through trust, anxiety, perceived risk, and the presence of human and institutional safeguards.

This research aims to address this gap by examining whether perceived AI usage competence is associated with AI acceptance in healthcare, whether this association is mediated by AI trust, and whether it is moderated by AI anxiety. Within this framework, the study assumes that perceived AI usage competence will increase trust, that trust will strengthen acceptance, and that usage competence will have both direct and indirect effects on acceptance through trust; conversely, anxiety towards AI will weaken both trust and acceptance. Thus, the study proposes an integrated model that considers both technical competence and psychosocial variables to explain AI acceptance in healthcare.

Literature Review and Hypothesis Development

The widespread adoption of AI applications in healthcare cannot be explained solely by algorithmic accuracy or technical capacity. The current literature shows that socio-cognitive variables, such as trust, perceived risk, explainability, data governance, user training, and experience, are decisive for AI acceptance in healthcare. In particular, a lack of knowledge and experience increases hesitation, uncertainty, and the perception of occupational risk among healthcare professionals. At the same time, appropriate training, transparency, and human supervision support a more balanced approach to AI. Therefore, the acceptance of AI in healthcare should be considered not only in terms of the user's level of exposure to technology, but also in conjunction with the extent to which they trust the technology, how competently they use it, and the emotional responses they have in this process (Asan et al., 2020; Choudhury & Shamszare, 2026; Steerling et al., 2023).

In this study, perceived AI usage competence refers to an individual's capacity to understand,

use, and interpret the outputs of AI tools; trust regarding AI refers to a positive expectation that the system will function consistently, usefully, and reliably; and AI acceptance in healthcare refers to a broader concept encompassing a positive attitude and tendency to use AI-assisted applications based on finding them appropriate, useful, and usable. Anxiety towards AI is a negative emotional response to perceived threats, such as making mistakes, loss of control, learning difficulties, ethical issues, data security, and job displacement. This conceptual framework supports considering AI trust as a mediating variable and AI anxiety as a moderating variable that may serve as a boundary condition (Asan et al., 2020; Shevtsova et al., 2024; Stevens & Stetson, 2023).

The Relationship Between Perceived AI Usage Competence, AI Trust, and AI Acceptance in Healthcare

In fields like healthcare, where the cost of error is high, trust is seen as a fundamental psychological mechanism for managing uncertainty. As users better understand how an AI system works, what types of inputs it relies on, under what conditions it provides reliable results, and when it requires human oversight, their evaluation of the technology becomes more well-founded. In other words, proficiency in use lays the groundwork for “calibrated” trust, rather than blind trust. Indeed, a lack of knowledge and a need for training in AI among healthcare professionals have been clearly reported; e-health literacy is positively correlated with technology trust, and trust increases with education/experience. Furthermore, systematic reviews highlight that the knowledge and experience gap regarding AI among healthcare professionals is a major obstacle to trust and adoption (Abdullah & Fakieh, 2020; Choudhury & Shamszare, 2026; Sahoo et al., 2025).

H₁. Perceived AI usage competence (X) is positively associated with AI trust (M).

In the technology acceptance literature, trust is one of the key variables that enables voluntary use

under uncertainty. This role is even more pronounced in healthcare, where users evaluate AI not only on technical performance but also on patient safety, accountability, professional autonomy, and ethical compliance. Stevens and Stetson (2023) showed that clinician trust accounted for 56% of the variance in acceptance for a particular AI application. Similarly, Kim, Choi, and Fotso (2024) reported that trust has a significant effect on acceptance intention for AI-based medical devices, and Kauttonen, Rousi, and Alamäki (2025) reported that trust and intention to use AI in healthcare are mutually shaped. Studies examining the stakeholder perspective also reveal that trust and acceptance are often fueled by the same determinants (Kauttonen et al., 2025; Shevtsova et al., 2024; Stevens & Stetson, 2023).

H₂. AI trust (M) is positively associated with AI acceptance in healthcare (Y).

The direct association between perceived AI usage competence and AI acceptance in healthcare can also be explained independently of AI trust. Competent individuals tend to view AI tools as more useful, more manageable, and more integrated into workflows because they better know how, when, and for what purpose to use them. Systematic review findings on healthcare professionals show that knowledge and experience gaps increase reservations about AI; conversely, education and skills development are critical for adoption. Empirically, e-health literacy among nurses has been shown to directly predict attitudes toward AI use, and productive AI literacy has been identified as a significant determinant of intention to use (Choe & Woo, 2025; Sahoo et al., 2025).

H₃. Perceived AI usage competence (X) is positively associated with AI acceptance in healthcare (Y).

The Mediating Role of AI Trust

The theory that user competence influences acceptance through trust is quite consistent. This is because competence helps the user perceive the AI system as more predictable, understandable,

and controllable; these cognitive evaluations translate into trust, which in turn fuels acceptance. Therefore, the link between competence and acceptance operates partly through a psychological trust mechanism rather than directly. This view is consistent with the Trust and Acceptance of Artificial Intelligence Technology (TrAAIT) model, which centers on clinician trust. Similarly, it has been shown that Unified Theory of Acceptance and Use of Technology (UTAUT) variables in medical AI devices can influence acceptance intention through trust, and that technology trust plays a mediating role in the relationship between e-health literacy and AI use attitudes (Chen et al., 2025; Kim et al., 2024; Stevens & Stetson, 2023).

H₄. AI trust (M) mediates the association between perceived AI usage competence (X) and AI acceptance in healthcare (Y).

The Moderator Role of AI Anxiety

AI anxiety reflects the extent to which an individual prioritizes potential threats related to the technology. When threats such as learning difficulties, erroneous decision-making, loss of functional control, data security, or job displacement become apparent, individuals may approach AI more cautiously despite their existing competence. In this case, even if user competence builds trust, anxiety can become a limiting factor that weakens this positive effect. Indeed, it has been shown that in primary nurses, AI learning anxiety and job replacement anxiety make sustained adoption difficult, while trust is a key psychological mechanism in this process. Furthermore, while practical AI experience reduces anxiety in Neonatal Intensive Care Units (NICU) nurses, awareness alone is not always reassuring; in nurses working in Türkiye, AI use, knowledge about AI, and finding AI trustworthy were found to be associated with lower anxiety and more positive attitudes (Ayed et al., 2025; Nirgiz et al., 2026; Zhou et al., 2025).

H₅. AI anxiety (W) moderates the direct effect of perceived AI usage competence (X) on AI trust (M).

AI anxiety may also attenuate the direct association between perceived AI usage competence

and AI acceptance in healthcare. Under normal circumstances, competence strengthens the assessment that the technology is useful and applicable; however, when anxiety increases, individuals may focus more on potential harm than benefit. Thus, the acceptance-generating effect of competence does not remain constant across all levels. This expectation is consistent with findings showing that AI anxiety in nurses is negatively correlated with both AI literacy and attitudes toward AI, and that anxiety functions as a significant psychological mechanism in the literacy-attitude link. Similarly, it has been reported that positive attitude scores decrease as anxiety about learning and AI configuration increases (Liu et al., 2025; Nirgiz et al., 2026; Zhou et al., 2025).

H₆. AI anxiety (W) moderates the direct effect of perceived AI usage competence (X) on AI acceptance in healthcare (Y).

Conditional Indirect Effect: The Role of AI Anxiety in the Relationship Between Perceived AI Usage Competence, Trust, and Acceptance

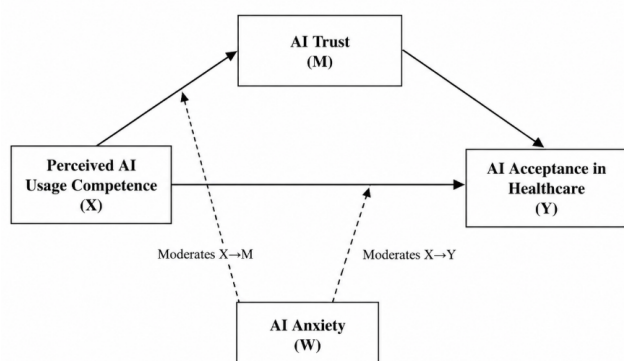
The conditional mediation hypothesis is a natural consequence of considering the above relationships together. If user competence increases trust, and trust increases acceptance, and if AI anxiety alters both the direct effect of user competence on trust and the direct effect of user competence on acceptance, then the trust-based indirect effect of user competence on acceptance is also expected to differ across levels of anxiety. In other words, at low anxiety levels, the reflection of user competence on trust and therefore on acceptance may be stronger. In contrast, at high anxiety levels, this effect may be weaker. Studies showing the mediating role of technology trust between literacy and attitude, and findings revealing the central role of AI anxiety in trust and sustained adoption processes, support this conditional mediation expectation (Chen et al., 2025; Liu et al., 2025; Zhou et al., 2025).

H₇. The indirect effect of perceived AI usage competence (X) on AI acceptance in healthcare (Y) varies across levels of AI anxiety (W).

Methodology

Research Model

This study is designed within the framework of a quantitative research approach, using a cross-sectional, correlational survey design. The theoretical model of the research is based on a moderated mediation model structure derived from Hayes' conditional process approach. In this model, perceived AI usage competence is considered the independent variable, AI trust is the mediating variable, AI acceptance in healthcare is the dependent variable, and AI anxiety is the moderating variable (Figure 1). According to the theoretical assumption, anxiety conditions both the effect of perceived AI usage competence on trust and the direct effect of perceived AI usage competence on AI acceptance in healthcare. Therefore, in this study, Hayes PROCESS Macro Model 8 was used in accordance with Hayes' regression-based conditional process approach (Hayes, 2022). As shown in Figure 1, AI anxiety was specified as a moderator of both the $X \rightarrow M$ path and the direct $X \rightarrow Y$ path, consistent with Hayes PROCESS Macro Model 8.



Note: X = Perceived AI Usage Competence; M = AI Trust; Y = AI Acceptance in Healthcare; W = AI Anxiety. AI anxiety moderates both the $X \rightarrow M$ and the direct $X \rightarrow Y$ paths, consistent with Hayes's PROCESS Model 8.

Figure 1. Proposed moderated mediation model of AI acceptance in healthcare

Sampling and Data Collection

The study's target population consisted of individuals aged 18 years and older who had received at least one healthcare service within the previous six

months. The accessible population consisted of adult individuals who accessed the research announcement during the online data collection and voluntarily agreed to participate in the study. Following ethical approval, data were collected between March 19, 2026, and April 1, 2026, via an online survey platform. The study used a combination of convenience sampling and snowball sampling, both non-probability sampling methods. The survey link was shared via email, social media, and messaging applications, and participants were also asked to share the link with other individuals who met the inclusion criteria.

Inclusion criteria for the study were being 18 years of age or older and having used healthcare services at least once in the last six months. The sample size was determined using a priori power analysis with G*Power 3.1 (Faul et al., 2009) to ensure sufficient statistical power for the regression-based analyses. Because Hayes PROCESS Model 8 is estimated through separate regression equations, the power analysis was conducted using a general multiple regression framework as an approximate criterion. In the focal model without covariates, the most complex outcome equation included four predictors: perceived AI usage competence, AI trust, AI anxiety, and the interaction term between perceived AI usage competence and AI anxiety. Under the assumptions of a moderate effect size ($f^2 = 0.15$), an $\alpha = .05$ significance level, and 95% statistical power, the minimum required sample size was 129 participants. The final analytic sample of 422 participants exceeded this minimum requirement. Therefore, this power analysis should be interpreted as an approximate assessment of sample adequacy for the regression-based model rather than as a direct power test for the conditional indirect association or the index of moderated mediation.

To account for potential data loss (e.g., incorrectly completed questionnaires) and to maintain sufficient analytical power after data cleaning, the target enrollment was approximately 400 participants. A total of 426 responses were initially recorded through the online survey platform. Before

the main analyses, the dataset was screened for eligibility, missing data, repetitive response patterns, and inconsistent responses to screening questions. Repetitive response patterns were evaluated by reviewing highly uniform responses across conceptually distinct scale items. Responses were excluded if participants did not meet the inclusion criteria, failed to provide electronic informed consent, or provided contradictory eligibility-related responses. In total, 4 responses were excluded during the data cleaning process, and the final analytic sample consisted of 422 participants.

Data Collection Tools

A structured online questionnaire was used as the data collection tool in the study. The questionnaire consists of four sections. These sections include electronic informed consent and screening questions, a personal information form prepared by the researcher, scales to measure the study's main variables, and contextual questions about healthcare use and AI experience.

Perceived AI usage competence was measured using all six items from the "Artificial Intelligence Knowledge and Usage Competence" subdimension of the Artificial Intelligence Literacy Scale developed by Aktay, Gök, and Kanata (2026). This subdimension was used because it directly reflects functional aspects of AI use, including understanding AI, selecting appropriate AI tools, adjusting AI-generated outputs according to the intended purpose, and providing effective prompts or commands. These competencies were considered relevant to the present study because the sample consisted of general adult healthcare users rather than technical experts, and the focus was on participants' perceived ability to understand and use AI-related tools in everyday and healthcare-related contexts. The use of the items was based on the published scale and its reported psychometric structure. No substantive changes were made to the conceptual content of the original items. Because the items were self-report statements, the construct was conceptualized as perceived AI usage competence rather than objectively assessed

AI competence. This choice is also consistent with the current international literature, which conceptualizes AI literacy along the dimensions of knowledge, usage, and evaluation (Aktay et al., 2026; Wang et al., 2023).

Trust in AI and acceptance of AI in healthcare were measured using the relevant sub-dimensions of the Social Perception of Artificial Intelligence in Healthcare Scale developed by Kuşcu Şahin (2026). The trust dimension comprises seven items, and the attitude and acceptance dimension comprises eight items. This scale addresses societal perception of AI in healthcare within a three-factor structure consisting of "attitudes and acceptance", "trust", and "perceived usefulness". It provides acceptable indicators of fit and reliability on an adult sample (Kuşcu Şahin, 2026).

Anxiety towards AI was assessed using seven items adapted from the AI Competence Anxiety Scale developed by Kaya, Güme, and Ergün (2026). This scale captures individuals' concerns about falling behind AI, feeling less competent in comparison with AI, and perceiving AI as a potential threat to their own abilities (Kaya et al., 2026). Although the scale was not developed specifically for healthcare AI contexts, it was considered relevant to the present study because the model focuses on general adult healthcare users' anxiety toward AI as a psychosocial boundary condition. In this context, competence-related AI anxiety may influence how individuals evaluate AI-supported healthcare applications, particularly when they feel uncertain about understanding, using, or interpreting AI-based systems. However, it should be noted that this measure may not fully capture healthcare-specific AI concerns such as diagnostic error, patient safety, data privacy, human supervision, clinical accountability, or the quality of human-patient interaction. Therefore, the findings related to AI anxiety were interpreted with this conceptual scope in mind.

All scale items were scored on a five-point Likert scale from 1 = "Strongly disagree" to 5 = "Strongly agree". A composite score was created by averaging the items for each dimension. Higher

scores indicate a higher level of the relevant construct.

Data Analysis

IBM SPSS Statistics and Hayes' PROCESS Macro add-on were used for data analysis. In the first stage of the analysis, the dataset was examined for missing data, outliers, multicollinearity, and dispersion characteristics. Categorical variables were summarized using frequencies and percentages, while continuous variables were summarized using mean, standard deviation, minimum, and maximum values. The internal consistency of the scales was evaluated using Cronbach's alpha, and pairwise relationships among the main variables were examined using Pearson's correlation. Common method bias was assessed using Harman's single-factor test because all variables were collected from the same participants through the same self-report questionnaire.

Hayes' PROCESS Macro Model 8 was used to test the hypothesized conditional process model. In this model, perceived AI usage competence was specified as the independent variable, AI trust as the mediating variable, AI acceptance in healthcare as the outcome variable, and AI anxiety as the moderating variable. Direct, indirect, conditional direct, and conditional indirect associations were estimated, and the index of moderated mediation was reported. Indirect associations were evaluated using 5,000 bootstrap samples and 95% bootstrap confidence intervals; confidence intervals that did not include zero were interpreted as statistically significant. To improve the interpretability of interaction terms, the independent variable and moderator were mean-centered before creating the interaction term, and conditional associations were examined at low, medium, and high levels of AI anxiety (Hayes, 2022). For interpretability, these conditional values were reported in the original scale metric. The hypothesized PROCESS Model 8 was first estimated without covariates to test the focal conditional process model. Subsequently, the model was rerun as a robustness check by including age, gender, education level, previous AI tool

use, healthcare use frequency, chronic disease status, and health insurance status as covariates.

Ethical Considerations

The research was deemed ethically appropriate by the Bayburt University Social and Human Sciences Research Ethics Committee with decision number 640 dated 18.03.2026. The research was conducted voluntarily, and participants' consent was explained in the questionnaire.

Participants were clearly informed on the first page of the survey about the purpose and scope of the study, that participation was voluntary, that responses would only be used for scientific purposes, that no identifying personal data would be collected, and that they could terminate the survey and withdraw from the study at any time. Only participants who provided electronic informed consent could access the survey questions.

No directly identifying information, such as name, identity number, telephone number, or full address, was collected in the study. A balanced research protocol between data security and data verification was followed, as protecting privacy and data integrity is critical in online data collection processes (Hardesty et al., 2024; Hohn et al., 2022).

Assumptions and Limitations

This research assumes that participants possess the cognitive ability to understand the questionnaire statements, that they answer the questions honestly and carefully, and that only individuals meeting the inclusion criteria participate in the questionnaire. Furthermore, it is assumed that the scales used will produce valid and reliable measurements within the research context.

The study has several limitations. First, the cross-sectional design of the research is not conducive to establishing temporal ordering among variables and making strong causal inferences. Second, the fact that data were collected online using self-report and non-probability sampling methods does not eliminate the risks of selection bias, content bias, social desirability bias, and common-method bias. Third, duplicate or fraudulent

submissions in anonymous or semi-anonymous online surveys can threaten data quality. Although data cleaning protocols were applied, it is not possible to eliminate this risk. Finally, since the study is limited to individuals aged 18 and over who have used health services in the last six months, generalization of the findings to different population groups should be done carefully (Huang & Shlomo, 2025; Savitz & Wellenius, 2023).

Results

The study included 422 participants, predominantly female, young adults (18-34 years), and bachelor's degree holders.

While more than half of the participants (56.2%) had previous experience with AI tools, a significant majority (66.8%) emphasized the necessity of human supervision in healthcare AI applications, citing data security (33.2%) and incorrect decisions (30.3%) as their primary concerns. Detailed demographic and background characteristics of the sample are presented in Table 1.

To test the scales' validity, a Confirmatory Factor Analysis (CFA) was conducted. Standardized factor loading ranges, model fit indices, Cronbach's alpha, composite reliability (CR), and average variance extracted (AVE) values were examined. The standardized factor loading ranges were reported.

Table 1. Descriptive Characteristics of Participants

Demographic / Individual Results		N	%
Gender	Female	233	55.2
	Male	189	44.8
Age	18-24	110	26.1
	25-34	128	30.3
	35-44	108	25.6
	45-54	53	12.6
	55 and above	23	5.5
Marital Status	Married	201	47.6
	Single	221	52.4
Education Level	Primary Education	21	5.0
	High school	82	19.4
	Associate Degree	108	25.6
	Bachelor's degree	149	35.3
	Postgraduate	62	14.7
Health Insurance	Yes	360	85.3
	No	62	14.7
Frequency of Healthcare Use in the Last 6 Months	1 time	157	37.2
	2-3 times	144	34.1
	4-5 times	79	18.7
	6 times or more	42	10.0
Chronic Disease	Yes	114	27.0
	No	308	73.0
Using AI Tools	Yes	237	56.2
	No	185	43.8
Human Supervision is Required	Yes	282	66.8
	No	63	14.9
	Undecided	77	18.2
The Biggest Concern	Wrong Decision	128	30.3
	Data Security	140	33.2
	Ethics	82	19.4
	No Concerns	72	17.1
Total		422	100.0

Table 2. Measurement Model, Reliability, and Convergent Validity Results

Fit Index	AI Usage Comp.	AI Accept. in Health.	AI Trust	AI Anxiety	Recommended Criterion*
SFLR*	.64-.73	.65-.75	.66-.74	.67-.72	≥0.50
χ^2 (df)	18.45(9)	33.09(20)	11.20(14)	25.58(14)	Reported
χ^2/df	2.05	1.65	0.80	1.83	<5
RMSEA	0.05	0.04	0.01	0.04	≤0.06
GFI	0.99	0.98	0.99	0.98	≥0.85
AGFI	0.97	0.97	0.98	0.97	≥0.80
CFI	0.99	0.99	1.00	0.99	≥0.95
NFI	0.98	0.98	0.99	0.98	≥0.95
IFI	0.99	0.99	1.00	0.99	≥0.95
TLI	0.98	0.99	1.00	0.98	≥0.95
SRMR	0.02	0.02	0.02	0.02	<0.08
AVE	0.51	0.50	0.52	0.50	≥0.50
CR	0.86	0.89	0.88	0.88	≥0.70
p	0.03	0.03	0.67	0.03	Reported
Cronbach's Alpha	0.839	0.889	0.864	0.865	

CR values above .70 were interpreted as evidence of acceptable construct reliability, and AVE values of .50 or higher were considered supportive of convergent validity. Discriminant validity was evaluated using the Fornell–Larcker criterion by comparing the square root of AVE values with inter-construct correlations (Fornell & Larcker, 1981; Hair et al., 2010). The square root of AVE for each construct was greater than its correlations with the other constructs, providing support for discriminant validity according to the Fornell–Larcker criterion (Table 3). The obtained goodness-of-fit values for the scales are given in Table 2. In addition to the individual CFA analyses, a combined four-factor measurement model including perceived AI usage competence, AI trust, AI acceptance in healthcare, and AI anxiety was tested. The combined model demonstrated acceptable model fit (χ^2 /df = 1.22, RMSEA = 0.02, GFI = 0.94, AGFI = 0.92, CFI = 0.99, NFI = 0.95, IFI = 0.99, TLI = 0.98, SRMR = 0.03, p = 0.003), supporting the distinctiveness of the four constructs. Given the conceptual proximity between AI trust and AI acceptance, HTMT analysis was additionally conducted. The HTMT value between AI trust and AI acceptance was 0.728 (below the recommended threshold of 0.90), providing further evidence that the two constructs are empirically distinguishable (Henseler et al., 2015).

The scales used in the study demonstrated high internal consistency.

Cronbach's alpha coefficients were calculated as 0.839 for perceived AI usage competence, 0.889 for AI acceptance in healthcare, 0.864 for AI trust, and 0.865 for AI anxiety (Table 2). These findings indicate that the measurements are internally consistent.

*Recommended criteria for model fit indices are based on Hu and Bentler (1999), Kline (2011), and Schermelleh-Engel et al. (2003). Criteria for standardized factor loadings, CR, and AVE are based on Fornell and Larcker (1981) and Hair et al. (2010). SFLR = standardized factor loading range; CR = composite reliability; AVE = average variance extracted.

Descriptive statistics indicated that participants' perceived AI usage competence and AI acceptance in healthcare were slightly above the midpoint, AI trust was at a moderate level, and AI anxiety was below the midpoint. Correlation analyses revealed significant positive relationships among perceived AI usage competence, AI trust, and AI acceptance. Conversely, AI anxiety was significantly and negatively correlated with all other variables in the model. Detailed descriptive statistics and correlation coefficients are provided in Table 3.

Common method bias was assessed using Harman's single-factor test because all study variables were collected from the same participants using the same questionnaire format. All measurement items were entered into an unrotated single-factor solution using principal component analysis. The

first factor accounted for 36.03% of the total variance, which was below the commonly used 50% threshold (Podsakoff et al., 2003). This result suggests that common method bias was unlikely to be a severe concern in the present dataset. However, given the cross-sectional self-report design, the possibility of common method variance cannot be completely ruled out.

tion factor (VIF) values were evaluated for all predictors. The tolerance values ranged from 0.543 to 0.996, and the VIF values ranged from 1.004 to 1.842, indicating that multicollinearity was not a serious concern. Because the independent variable and moderator were mean-centered before creating the interaction term, nonessential multicollinearity associated with the product term was also reduced.

Table 3. Descriptive Statistics, Correlations, and Discriminant Validity

SCALES	M	SD	Skew.	Kurt.	1	2	3	4
1. AI Usage Comp. (X)	3.23	0.72	-0.140	-0.702	(0.71)			
2. AI Accept. in Healthcare (Y)	3.30	0.68	0.253	-0.253	0.550*	(0.71)		
3. AI Trust (M)	3.10	0.67	0.510	0.178	0.608*	0.642*	(0.72)	
4. AI Anxiety (W)	2.77	0.69	0.122	-0.269	-0.378*	-0.490*	-0.441*	(0.71)

Note: Values in parentheses on the diagonal represent the square root of AVE. Values below the diagonal represent Pearson correlation coefficients.

* $p < 0.01$

To test the research hypotheses, mediation analysis was performed using the Hayes PROCESS Macro within the framework of Model 8. In this model, perceived AI usage competence was the independent variable (X), AI trust the mediating variable (M), AI acceptance in healthcare the dependent variable (Y), and AI anxiety the moderating variable (W).

To avoid multicollinearity in moderating effects analysis, the correlation between the independent (perceived AI Usage Competence) and moderating (AI Anxiety) variables should not exceed 0.90 (Hair et al., 2010). The correlation between perceived AI usage competence and AI anxiety ($r = -0.378$, $p < 0.01$) was no greater than 0.90, indicating the absence of multicollinearity. Furthermore, tolerance and VIF values were examined to assess multicollinearity. Tolerance values above 0.10 or VIF values below 10 indicate the absence of multicollinearity (Pallant, 2007). In the mediator equation, perceived AI usage competence, AI anxiety, and the interaction term between perceived AI usage competence and AI anxiety were examined. In the outcome equation, perceived AI usage competence, AI trust, AI anxiety, and the interaction term between perceived AI usage competence and AI anxiety were examined. Tolerance and variance infla-

According to the mediating-variable model results, perceived AI usage competence was positively and significantly associated with AI trust ($b = 0.493$; $SE = 0.036$; $t = 13.547$; $p < 0.001$). This finding supports hypothesis H₁. AI anxiety was negatively and significantly associated with trust is negative and significant ($b = -0.236$; $SE = 0.038$; $t = -6.261$; $p < 0.001$). Furthermore, the interaction term between perceived AI usage competence and AI anxiety was significant ($b = -0.237$; $SE = 0.046$; $t = -5.184$; $p < 0.001$). This result shows that as AI anxiety is associated with higher levels of the positive effect of perceived AI usage competence on trust attenuates the association, supporting hypothesis H₅. It was also determined that the interaction term added to the model made a significant contribution to the model ($\Delta R^2 = 0.035$; $F(1,418) = 26.878$; $p < 0.001$). The explanatory level of the mediating variable model is 45.7% ($R^2 = 0.457$; $F(3,418) = 117.276$; $p < 0.001$).

According to the dependent variable model results, AI trust was positively and significantly associated with AI acceptance in healthcare ($b = 0.394$; $SE = 0.048$; $t = 8.185$; $p < 0.001$). This finding supports hypothesis H₂. The direct association between perceived AI usage competence and AI acceptance in healthcare was also found to be positive and significant ($b = 0.222$; $SE = 0.043$; $t = 5.178$; $p < 0.001$), thus supporting hypothesis H₃. On the other hand, AI anxiety was negatively and significantly associated with AI acceptance in healthcare ($b = -0.231$; $SE =$

0.039; $t = -5.961$; $p < 0.001$). The interaction term between perceived AI usage competence and AI anxiety was statistically significant in the acceptance model ($b = -0.103$; $SE = 0.046$; $t = -2.213$; $p = 0.027$), suggesting that the direct association between perceived AI usage competence and AI acceptance in healthcare varied across levels of AI anxiety and supports hypothesis H₆. However, the increase in explained variance associated with this interaction was small ($\Delta R^2 = 0.006$; $F(1,417) = 4.898$; $p < 0.027$); therefore, the practical magnitude of this moderation finding should be interpreted cautiously. The explanatory power of the outcome model was 49.9% ($R^2 = 0.499$; $F(4,417) = 103.86$; $p < 0.001$).

This conditional association was $b = 0.656$ at the low anxiety level and decreased to $b = 0.329$ at the high anxiety level. Similarly, the direct effect of perceived AI usage competence on AI acceptance in healthcare was stronger at the low anxiety level ($b = 0.293$) and weaker at the high anxiety level ($b = 0.152$).

Bootstrap results indicate that the trust-mediated indirect association of perceived AI usage competence on AI acceptance in healthcare was statistically significant at all examined levels of AI anxiety, as none of the bootstrap confidence intervals included zero.

Table 4. Regression Results for Hayes PROCESS Model 8

Model	Predictor	b	SE	t	p	95% CI
Trust Model	Perc. AI Usage Competence (X)	0.493	0.036	13.547	< .001	[0.421, 0.564]
	AI Anxiety (W)	-0.236	0.038	-6.261	< .001	[-0.310, -0.162]
	Interaction Term (X × W)	-0.237	0.046	-5.184	< .001	[-0.327, -0.147]
	$R^2 = .457$, $F(3, 418) = 117.276$, $p < .001$; $\Delta R^2 = .035$, $F(1, 418) = 26.878$, $p < .001$					
Acceptance Model	Perc. AI Usage Competence (X)	0.222	0.043	5.178	< .001	[0.138, 0.307]
	AI Trust (M)	0.394	0.048	8.185	< .001	[0.299, 0.489]
	AI Anxiety (W)	-0.231	0.039	-5.962	< .001	[-0.308, -0.155]
	Interaction Term (X × W)	-0.103	0.046	-2.213	.027	[-0.194, -0.011]
$R^2 = .499$, $F(4, 417) = 103.864$, $p < .001$; $\Delta R^2 = .006$, $F(1, 417) = 4.898$, $p = .027$						

To examine the moderating role of AI anxiety, conditional associations were examined at one standard deviation below the mean, the mean, and one standard deviation above the mean of AI anxiety. The findings indicate that the association between perceived AI usage competence and AI trust is positive and statistically significant at all levels of AI anxiety; however, this association becomes weaker as AI anxiety increases.

Because both the X→M path and the direct X→Y path were modeled as conditional on AI anxiety in PROCESS Model 8, the indirect and total associations were interpreted conditionally rather than as single unconditional estimates.

At the mean level of AI anxiety, the trust-mediated indirect association was statistically significant ($b = 0.194$, Boot SE = 0.026, Boot 95% CI [0.146, 0.245]), supporting the mediating role of AI trust.

Table 5. Conditional Direct Effects, Conditional Indirect Effects, and Moderated Mediation

AI Anxiety (W)	X→M Effect (b [95% CI])	X→Y Direct Effect (b [95% CI])	X→M→Y Indirect Effect b [95% CI])
Low (-1 SD = 2.08)	0.656 [0.559, 0.754]	0.293 [0.179, 0.408]	0.259 [0.194, 0.328]
Medium (Mean = 2.77)	0.493 [0.421, 0.564]	0.222 [0.138, 0.307]	0.194 [0.146, 0.245]
High (+1 SD = 3.46)	0.329 [0.238, 0.421]	0.152 [0.056, 0.247]	0.130 [0.088, 0.176]
Index of Moderated Mediation = -0.094 [-0.136, -0.057]			

Note: X=Perc. AI Usage Competence; M=AI Trust; Y=AI Acceptance in Healthcare. The independent variable and moderator were mean-centered before creating the interaction term. Conditional associations were estimated at -1 SD, the mean, and +1 SD of AI anxiety. Bootstrap confidence intervals are based on 5,000 samples.

The conditional direct association at the mean level of AI anxiety was also statistically significant ($b = 0.222$, $SE = 0.043$, 95% CI $[0.138, 0.307]$); therefore, the conditional total association at the mean level of AI anxiety was 0.417. However, the magnitude of the indirect effect decreased as anxiety increased. The significant finding of the index of moderated mediation (index = -0.094 ; Boot SE = 0.020 ; Boot 95% CI $[-0.136, -0.057]$) indicates that the indirect effect of perceived AI usage competence on AI acceptance in healthcare varies across levels of AI anxiety. These results support hypotheses H_4 and H_7 .

As a covariate-adjusted robustness check, Hayes PROCESS Model 8 was rerun by including age, gender, education level, previous AI tool use, healthcare use frequency, chronic disease status, and health insurance status as covariates. The covariate-adjusted model produced a substantively similar pattern of findings. Perceived AI usage competence remained positively associated with AI trust ($b = 0.449$, $SE = 0.045$, $p < .001$, 95% CI $[0.360, 0.537]$) and AI acceptance in healthcare ($b = 0.188$, $SE = 0.050$, $p < .001$, 95% CI $[0.090, 0.287]$). AI trust remained positively associated with AI acceptance in healthcare ($b = 0.402$, $SE = 0.049$, $p < .001$, 95% CI $[0.305, 0.499]$), and AI anxiety remained negatively associated with both AI trust ($b = -0.270$, $SE = 0.040$, $p < .001$, 95% CI $[-0.348, -0.193]$) and AI acceptance in healthcare ($b = -0.210$, $SE = 0.042$, $p < .001$, 95% CI $[-0.292, -0.129]$). The interaction terms also remained statistically significant in both the trust model ($b = -0.247$, $SE = 0.046$, $p < .001$, 95% CI $[-0.336, -0.157]$) and the acceptance model ($b = -0.107$, $SE = 0.047$, $p = .024$, 95% CI $[-0.199, -0.014]$). The index of moderated mediation remained significant (index = -0.099 , Boot SE = 0.021 , Boot 95% CI $[-0.141, -0.060]$), supporting the robustness of the moderated mediation findings after adjustment for covariates.

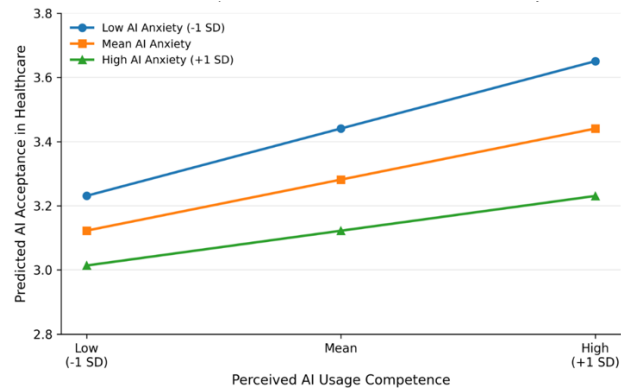


Figure 2. Conditional Direct Effect

Figure 2 illustrates the conditional direct association between perceived AI usage competence and AI acceptance in healthcare at low, mean, and high levels of AI anxiety. The slope was steepest at low AI anxiety and flattest at high AI anxiety, indicating that the positive association between perceived AI usage competence and AI acceptance in healthcare was weaker among participants with higher AI anxiety.

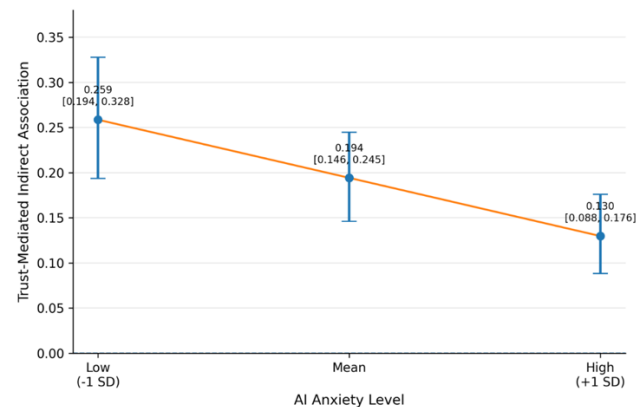


Figure 3. Conditional Indirect Effects and Confidence Intervals

Figure 3 visualizes how the indirect effect of perceived AI usage competence on AI acceptance via the trust (M) mechanism ($X \rightarrow M \rightarrow Y$) varies across different levels (Low, Medium, High) of AI anxiety (W), the moderating variable. The vertical lines (error bars) represent the 95% Bootstrap confidence intervals (Boot LLCI and Boot ULCI) calculated for each anxiety level. The absence of a zero (0) value in the confidence intervals at all anxiety levels indicates that the trust-mediated indirect association was statistically significant under all conditions. The overall downward trend in the

graph suggests that the trust-mediated indirect association between perceived AI usage competence and AI acceptance in healthcare gradually decreases as individuals' AI anxiety increases.

Taken together, the findings indicate that perceived AI usage competence was positively associated with AI trust and AI acceptance in healthcare, and that AI trust was positively associated with AI acceptance. The trust-mediated indirect association was statistically significant, and AI anxiety attenuated both the competence–trust association and the direct competence–acceptance association. The index of moderated mediation further indicated that the trust-mediated indirect association varied across levels of AI anxiety. Accordingly, all hypotheses H₁, H₂, H₃, H₄, H₅, H₆, and H₇ are supported.

Discussion and Conclusion

This study suggests that perceived AI usage competence is positively associated with AI trust and AI acceptance in healthcare, and that AI trust may help explain the association between perceived competence and acceptance. The findings also indicate that AI anxiety functions as a boundary condition, as the competence–trust and competence–acceptance associations were weaker among participants reporting higher levels of AI anxiety. Accordingly, the proposed moderated mediation model should be interpreted as a theoretically specified pattern of associations among the study variables rather than evidence of temporal ordering or causal mechanisms.

The positive associations among perceived AI usage competence, trust, and acceptance suggest that users who perceive themselves as more competent in understanding and using AI may evaluate healthcare AI more favorably. This interpretation is consistent with studies showing that AI literacy, knowledge, and experience are related to more positive attitudes toward AI and to greater acceptance of healthcare AI applications (Çapuk et al., 2026; Kauttonen et al., 2025; Sahoo et al., 2025). However, because perceived AI usage competence was measured through self-report items,

the findings reflect perceived rather than objectively assessed competence.

The findings also highlight the central role of AI trust in healthcare AI acceptance. In healthcare, trust may be especially important because users evaluate AI systems not only in terms of usefulness, but also in relation to privacy, accountability, transparency, ethical compliance, and human supervision. This is consistent with previous healthcare AI literature emphasizing the role of explainability, governance, perceived value, and contextual safeguards in shaping trust and acceptance (Asan et al., 2020; Shevtsova et al., 2024; Zhou et al., 2025).

AI anxiety was found to attenuate the positive associations between perceived competence, trust, and acceptance. This suggests that even when users perceive themselves as competent, concerns about errors, loss of control, data privacy, ethical uncertainty, or reduced human involvement may limit the extent to which perceived competence is associated with trust and acceptance. This finding is consistent with previous studies showing that AI-related anxiety is negatively associated with favorable attitudes and sustained use intentions (Nirgiz et al., 2026; Zhou et al., 2025).

The descriptive findings also suggest that healthcare users may hold a conditional form of acceptance toward AI in healthcare. Many participants emphasized the need for human supervision and reported concerns about data security and incorrect decisions. This pattern is consistent with studies indicating that trust in healthcare AI may increase when AI systems are supported by provider oversight, regulatory approval, transparent performance information, and clear safeguards (Erul et al., 2025; Kirchberger, 2026; Zhu et al., 2026).

The study contributes to the healthcare AI acceptance literature by examining perceived AI usage competence, AI trust, and AI anxiety within the same conditional process model among adult healthcare users in Türkiye. Rather than substantially extending or replacing traditional technology acceptance models, the findings provide complementary evidence that psychosocial factors such

as trust and anxiety should be considered alongside perceived competence when examining AI acceptance in healthcare (Scipion et al., 2025; Stevens & Stetson, 2023).

From a practical perspective, the findings suggest that communication strategies for patient-facing healthcare AI applications should not focus only on technical functionality. Clear information about algorithmic limitations, potential sources of error, data privacy, ethical responsibilities, accountability mechanisms, and the role of healthcare professionals may be important for supporting trust and reducing anxiety among healthcare users (Asan et al., 2020; Lekadir et al., 2025). Transparent governance mechanisms that prioritize accountability, data protection, and responsible implementation may also support more favorable conditions for trust in healthcare AI (Hasan et al., 2024; Henzler et al., 2025; Kauttonen et al., 2025). However, because the sample consisted of general adult healthcare users rather than healthcare administrators or clinicians, organizational implications should be interpreted cautiously.

Several limitations should be considered. First, the cross-sectional design does not allow conclusions about temporal ordering or causality. Second, the use of online convenience and snowball sampling limits generalizability. Third, all variables were measured using self-report scales at a single time point. Although Harman's single-factor test did not indicate severe common method bias, common method variance cannot be completely ruled out. Finally, AI anxiety was measured using a competence-related AI anxiety scale rather than a healthcare-specific AI anxiety measure; therefore, it may not fully capture concerns such as diagnostic error, patient safety, data privacy, human supervision, and clinical accountability.

Future research should use longitudinal, experimental, and scenario-based designs to examine how perceived competence, trust, anxiety, explainability, performance information, human supervision, and governance mechanisms shape AI acceptance across different healthcare AI use cases

(Scipion et al., 2025; Zhu et al., 2026). Future studies should also include more diverse samples and compare different stakeholder groups, including patients, clinicians, healthcare managers, and policymakers. Overall, the findings suggest that AI acceptance in healthcare is associated not only with perceived technical competence, but also with trust, anxiety, and perceived safeguards.

Declarations

Funding: The author received no financial support for the research, authorship, and/or publication of this article.

Conflicts of Interest: The author declares no conflicts of interest related to this study.

Ethical Approval: Ethical committee approval for this study was obtained from Social and Human Sciences Research Ethics Committee of Bayburt University on March 18, 2026 (decision/document number: E-51694156-050.04-355402).

Informed Consent: Electronic informed consent was obtained from all participants included in the study (Individuals aged 18 and over who have received healthcare services at least once within the last six months).

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

AI Disclosure: *The scholarly content of this work was entirely produced by the author. AI-based tools, namely ChatGPT, Gemini, and Grammarly, were used only for limited language editing, fluency improvement, and formal checks. These tools were not used for data collection or generation, statistical analysis, or citation/reference generation. The author reviewed, edited, and verified all AI-assisted outputs and takes full responsibility for the content of the manuscript.*

References

- Abdullah, R., & Fakieh, B. (2020). Health care employees' perceptions of the use of artificial intelligence applications: Survey study. *Journal of Medical Internet Research*, 22(5), e17620. <https://doi.org/10.2196/17620>
- Aktay, S., Gök, S., & Kanata, O. (2026). Artificial intelligence literacy scale. *International Technology and Education Journal*, 10(1). <https://doi.org/10.71410/ESApub.ITEJ.2026.10-01.02>
- Asan, O., Bayrak, A. E., & Choudhury, A. (2020). Artificial intelligence and human trust in healthcare: Focus on clinicians. *Journal of Medical Internet Research*, 22(6), e15154. <https://doi.org/10.2196/15154>
- Ayed, A., Batran, A., Aqtam, I., Malak, M. Z., Abu Ejheisheh, M., Farajallah, M., Farraj, L., & Alkhatib, S. (2025). Perceived worries in the adoption of artificial intelligence among nurses in neonatal intensive care units. *BMC Nursing*, 24, 777. <https://doi.org/10.1186/s12912-025-03318-z>
- Bracic, A., Spector-Bagdady, K., Towle, S., Zhang, R., James, C. A., & Price, W. N. (2026). Factors for patient trust and acceptance of medical artificial intelligence. *JAMA Network Open*, 9(3), e260815. <https://doi.org/10.1001/jamanetworkopen.2026.0815>
- Chen, Z., Wu, T., Wu, X., Wang, J., Ke, S., Li, H., & Lin, R. (2025). The mediating effects of technology trust and perceived value in the relationship between eHealth literacy and attitude toward the usage of artificial intelligence in nursing: A cross-sectional study. *BMC Nursing*, 24(1), 989. <https://doi.org/10.1186/s12912-025-03577-w>
- Choe, J., & Woo, K. (2025). Factors associated with intention to use generative artificial intelligence in nursing practice: A cross-sectional study. *BMC Nursing*, 24(1), 1327. <https://doi.org/10.1186/s12912-025-03985-y>
- Choudhury, A., & Shamszare, H. (2026). Human factors influencing trust in healthcare artificial intelligence: Systematic literature review. *IJSE Transactions on Occupational Ergonomics and Human Factors*. Advance online publication. <https://doi.org/10.1080/24725838.2026.2625655>
- Çapuk, H., Yiğit, M. F., & Uçar, M. (2026). Exploring the relationship between health professionals' artificial intelligence literacy and their attitudes toward artificial intelligence. *Informatics for Health and Social Care*. Advance online publication. <https://doi.org/10.1080/17538157.2025.2602515>
- Erul, E., Aktekin, Y., Danişman, F. B., Gümüştaş, Ş. A., Saraç Aktekin, B., Yekedüz, E., & Ürün, Y. (2025). Perceptions, attitudes, and concerns on artificial intelligence applications in patients with cancer. *Cancer Control*, 32, 10732748251343245. <https://doi.org/10.1177/10732748251343245>
- Faul, F., Erdfelder, E., Buchner, A., & Lang, A.-G. (2009). Statistical power analyses using G*Power 3.1: Tests for correlation and regression analyses. *Behavior Research Methods*, 41(4), 1149–1160. <https://doi.org/10.3758/BRM.41.4.1149>
- Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39–50.
- Hair, J. F., Jr., Black, W. C., Babin, B. J., & Anderson, R. E. (2010). *Multivariate data analysis* (7th ed.). Prentice Hall.
- Hardesty, J. J., Crespi, E., Sinamo, J. K., Nian, Q., Breland, A., Eissenberg, T., Kennedy, R. D., & Cohen, J. E. (2024). From doubt to confidence—Overcoming fraudulent submissions by bots and other takers of a web-based survey. *Journal of Medical Internet Research*, 26, e60184. <https://doi.org/10.2196/60184>
- Hassan, M., Kushniruk, A., & Borycki, E. (2024). Barriers to and facilitators of artificial intelligence adoption in health care: Scoping review. *JMIR Human Factors*, 11, e48633. <https://doi.org/10.2196/48633>
- Hayes, A. F. (2022). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach* (3rd ed.). Guilford Press.
- 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>
- Henzler, D., Schmidt, S., Koçar, A., Herdegen, S., Lindinger, G. L., Maris, M. T., Bak, M. A. R., Willems, D. L., Tan, H. L., Lauerer, M., Nagel, E., Hindricks, G., Dagres, N., & Konopka, M. J. (2025). Healthcare professionals' perspectives on artificial intelligence in patient care: A systematic review of hindering and facilitating factors on different levels. *BMC Health Services Research*, 25(1), 633. <https://doi.org/10.1186/s12913-025-12664-2>
- Hohn, K. L., Braswell, A. A., & DeVita, J. M. (2022). Preventing and protecting against internet research fraud in anonymous web-based research: Protocol for the development and implementation of an anonymous web-based data integrity plan. *JMIR Research Protocols*, 11(9), e38550. <https://doi.org/10.2196/38550>
- Hu, L. T., & Bentler, P. M. (1999). Cutoff criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives. *Structural Equation Modeling: A Multidisciplinary Journal*, 6(1), 1–55. <https://doi.org/10.1080/10705519909540118>
- Huang, Y., & Shlomo, N. (2025). Improving the reliability of a nonprobability web survey: Application to measuring gender pay gap. *International Journal of Social Research Methodology*, 28(3), 277–290. <https://doi.org/10.1080/13645579.2024.2371238>
- Kauttonen, J., Rousi, R., & Alamäki, A. (2025). Trust and acceptance challenges in the adoption of AI applications in health care: Quantitative survey analysis. *Journal of Medical Internet Research*, 27, e65567. <https://doi.org/10.2196/65567>
- Kaya, B., Güme, S., & Ergün, N. (2026). The age of human comparison with AI: Development of the AI Competence Anxiety Scale and the AI Optimism Scale, and their relationship with generative AI literacy and AI self-efficacy. *International Journal of Human-Computer Interaction*. Advance online publication. <https://doi.org/10.1080/10447318.2025.2609913>
- Kim, Y. J., Choi, J. H., & Fotso, G. M. N. (2024). Medical professionals' adoption of AI-based medical devices: UTAUT model with trust mediation. *Journal of Open Innovation: Technology, Market, and Complexity*, 10(1), 100220. <https://doi.org/10.1016/j.joitmc.2024.100220>
- Kirchberger, M. (2026). Public perceptions of artificial intelligence in medicine and implications for future medical education: A cross-sectional survey. *JMIR Formative Research*. Advance online publication. <https://doi.org/10.2196/89123>
- Kline, R. B. (2011). *Principles and practice of structural equation modeling* (3rd ed.). Guilford Press.
- Kuşcu Şahin, F. N. (2026). Development and examination of the psychometric properties of the social perception of artificial intelligence in healthcare scale in the Turkish context: Evidence from Hatay Province. *International Journal of Public Health*, 71, 1609194. <https://doi.org/10.3389/ijph.2026.1609194>
- LaGore, A., Colby, L., Sinsabaugh, K., Grobbel, C., & Piscotty, R. (2025). The effect of AI education on nursing student AI literacy and anxiety. *Journal of Nursing Education*, 64(9), 587–590. <https://doi.org/10.3928/01484834-20250415-04>
- Lekadir, K., Frangi, A. F., Porras, A. R., Glocker, B., Cintas, C., Langlotz, C. P., et al.; FUTURE-AI Consortium. (2025). FUTURE-AI: International consensus guideline for trustworthy and deployable artificial intelligence in healthcare. *BMJ*, 388, e081554. <https://doi.org/10.1136/bmj-2024-081554>
- Li, W., & Liu, X. (2025). Anxiety about artificial intelligence from patient and doctor-physician. *Patient Education and Counseling*, 133, 108619. <https://doi.org/10.1016/j.pec.2024.108619>
- Liu, S., Xiao, Y., Nie, M., Yuan, X., Wang, L., Wang, M., & Wu, Z. (2025). Nurses' attitudes toward artificial intelligence: AI literacy as a predictor and the mediating effect of AI anxiety. *BMC Nursing*, 24(1), 1511. <https://doi.org/10.1186/s12912-025-04142-1>
- Nirgiz, C., Kıymaç Sarı, M., & Çaylı, N. (2026). The relationship between nurses' anxiety and attitudes towards artificial intelligence and examination of influencing factors. *BMC Nursing*, 25(1), 122. <https://doi.org/10.1186/s12912-026-04293-9>

- Pallant, J. (2007). *SPSS survival manual* (3rd ed.). Open University Press.
- Podsakoff, P. M., MacKenzie, S. B., Lee, J.-Y., & Podsakoff, N. P. (2003). Common method biases in behavioral research: A critical review of the literature and recommended remedies. *Journal of Applied Psychology*, 88(5), 879–903. <https://doi.org/10.1037/0021-9010.88.5.879>
- Sahoo, R. K., Sahoo, K. C., Negi, S., Baliarsingh, S. K., Panda, B., & Pati, S. (2025). Health professionals' perspectives on the use of artificial intelligence in healthcare: A systematic review. *Patient Education and Counseling*, 134, 108680. <https://doi.org/10.1016/j.pec.2025.108680>
- Savitz, D. A., & Wellenius, G. A. (2023). Can cross-sectional studies contribute to causal inference? It depends. *American Journal of Epidemiology*, 192(4), 514–516. <https://doi.org/10.1093/aje/kwac037>
- Schermelleh-Engel, K., Moosbrugger, H., & Müller, H. (2003). Evaluating the fit of structural equation models: Tests of significance and descriptive goodness-of-fit measures. *Methods of Psychological Research*, 8, 23–74.
- Scipion, C. E. A., Manchester, M. A., Federman, A., Wang, Y., & Arias, J. J. (2025). Barriers to and facilitators of clinician acceptance and use of artificial intelligence in healthcare settings: A scoping review. *BMJ Open*, 15(4), e092624. <https://doi.org/10.1136/bmjopen-2024-092624>
- Shevtsova, D., Ahmed, A., Boot, I. W. A., Sanges, C., Hudecek, M., Jacobs, J. J. L., Hort, S., & Vrijhoef, H. J. M. (2024). Trust in and acceptance of artificial intelligence applications in medicine: Mixed methods study. *JMIR Human Factors*, 11, e47031. <https://doi.org/10.2196/47031>
- Steerling, E., Siira, E., Nilsen, P., Svedberg, P., & Nygren, J. (2023). Implementing AI in healthcare—the relevance of trust: A scoping review. *Frontiers in Health Services*, 3, 1211150. <https://doi.org/10.3389/frhs.2023.1211150>
- Stevens, A. F., & Stetson, P. (2023). Theory of trust and acceptance of artificial intelligence technology (TrAAIT): An instrument to assess clinician trust and acceptance of artificial intelligence. *Journal of Biomedical Informatics*, 148, 104550. <https://doi.org/10.1016/j.jbi.2023.104550>
- Švestková, A., Huang, Y., & Šmahel, D. (2025). Factors that influence trust and willingness to use generative AI for health information: A cross-sectional study. *Digital Health*, 11, 20552076251360973. <https://doi.org/10.1177/20552076251360973>
- Wang, B., Rau, P. L. P., & Yuan, T. (2023). Measuring user competence in using artificial intelligence: Validity and reliability of artificial intelligence literacy scale. *Behaviour & Information Technology*, 42(9), 1324–1337. <https://doi.org/10.1080/0144929X.2022.2072768>
- Zhou, Q., Yang, L., Tang, Y., Yang, J., Zhou, W., Guan, W., Yan, L., & Liu, Y. (2025). The mediation of trust on artificial intelligence anxiety and continuous adoption of artificial intelligence technology among primary nurses: A cross-sectional study. *BMC Nursing*, 24(1), 724. <https://doi.org/10.1186/s12912-025-03406-0>
- Zhu, X., Stroud, A. M., Minter, S. A., Yoo, D. W., Ridgeway, J. L., Mooghali, M., Miller, J. E., & Barry, B. A. (2026). Key information influencing patient decision-making about AI in health care: Survey experiment study. *Journal of Medical Internet Research*, 28, e75615. <https://doi.org/10.2196/75615>