

Dual Trajectories of User Engagement with Medical Generative AI: Cognitive Maturation and Attitudinal Recalibration^{*}

Yibo Meng^{1,*}, Qiuyu Long^{2,†}, Bingyi Liu³, Yan Guan⁴ and Xiaolan Ding⁵

¹Cornell University, United States

²Department of Big Data Management and Application, Nanjing University of Chinese Medicine, China

³University of Michigan, Ann Arbor, United States

⁴Tsinghua University, China

⁵Zhengzhou Zhongmu County First People's Hospital, China

Abstract

This paper presents a four-year longitudinal mixed-methods study conducted from 2022 to 2025. It investigates the evolving attitudes of users toward conversational and generative AI-based Intelligent User Interfaces (IUIs) in healthcare. The study initiated with 21 participants and retained 16 through completion. Findings reveal a dual-trajectory model of engagement. While participant AI literacy demonstrated a consistent longitudinal increase, overall support followed a non-linear trajectory. This path involved a phase of optimistic acceptance peaking in 2024 followed by a period of discerning recalibration in 2025. Quantitative analysis using RM-ANOVA confirmed aspects of this pattern. Specifically, Perceived Usefulness showed a significant post-peak decrease ($p = .006$), while behavioral acceptance stabilized. Qualitative data revealed that this recalibration was driven by deepening critical awareness of ethical complexities including privacy, accountability, and algorithmic fairness rather than technical dissatisfaction. This work contributes a Critical-Dynamic Framework to HCI. It challenges linear technology acceptance models by demonstrating how cognitive maturation triggers complex trust recalibration. We translate these insights into design principles for next-generation medical IUIs. We advocate for systems that prioritize trust calibration over maximization, foster human-AI collaboration over replacement, and embed ethics and adaptive transparency as core infrastructure. These principles aim to support sustainable long-term adoption in high-stakes healthcare contexts.

Keywords

Medical Generative AI, Trust Calibration, Longitudinal Study, Intelligent User Interfaces, Attitudinal Recalibration, Human-AI Collaboration

1. Introduction

The emergence of large-scale generative artificial intelligence (AI) has initiated a significant transformation in human-computer interaction, and the healthcare domain faces substantial potential for disruption[1][2]. Intelligent User Interfaces (IUIs) powered by generative AI are transitioning from theoretical concepts to tangible applications. These applications manifest as AI health chatbots, mental wellness systems, and personalized health assistants[3][4][5]. Such systems aim to democratize access to medical information, provide scalable support for chronic disease management, and offer personalized guidance for daily well-being[6][7]. By simulating empathetic conversation and generating tailored health plans, these IUIs attempt to function as integral partners in the health management process of the user[4][8].

However, the effective deployment of such systems in the high-stakes domain of healthcare presents profound Human-Computer Interaction (HCI) challenges that extend beyond technological capability[9][10]. The relationship between a user and a medical IUI constitutes an evolving partner-

Joint Proceedings of the ACM Intelligent User Interfaces (IUI) Workshops 2026, March 23-26, 2026, Paphos, Cyprus

^{*}Corresponding author.

[†]These authors contributed equally.

✉ mengyb22@tsinghua.org.cn (Y. Meng); 012922139@njucm.edu.cn (Q. Long); bingyi@umich.edu (B. Liu); guany@tsinghua.edu.cn (Y. Guan); 13838038961@163.com (X. Ding)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

ship rather than a single transactional event, and it relies on the dynamic maintenance of trust[11]. The success of these technologies depends on the ability of the interface to facilitate appropriate trust calibration. This process involves guiding users to understand the strengths of the system as well as its limitations and uncertainties. While initial user acceptance is necessary, the long-term trajectory of this trust determines whether these systems are sustainably adopted or ultimately abandoned. This trajectory is shaped by ongoing interactions, evolving user literacy, and external technological developments. The IUI serves as the primary locus of this negotiation, as it mediates the mental model the user holds regarding the complex internal mechanisms of the AI[12].

Existing research on technology acceptance in health has predominantly relied on cross-sectional studies or static models. These approaches often fail to capture the dynamic and co-adaptive nature of the human-AI relationship[13][14]. User perceptions are continuously reshaped as users gain experience and move from novice to expert status. Furthermore, the technology itself evolves rapidly in the public consciousness[8]. A user who initially approaches a health AI with simple queries may learn to probe its reasoning, challenge its outputs, and use it as a sophisticated tool for preparing for clinical encounters over time[15][16]. This progression from initial curiosity to instrumentalization and potentially to critical re-evaluation represents a gap in current understanding. A longitudinal perspective is therefore essential to uncover the mechanisms that govern the long-term viability and ethical integration of IUIs in healthcare.

To address this gap, this paper presents findings from a four-year longitudinal, mixed-methods study (2022–2025). The study tracks the evolving attitudes of an initial cohort of 21 participants, with 16 completing the study, toward conversational and generative AI in medicine. This period spans the technological transition from earlier conversational agents to advanced large language models. We chart the complex interplay between growing AI literacy among users and their shifting sentiments to reveal a nuanced narrative of technological domestication. Our findings indicate a divergence. While AI literacy and practical skills demonstrated a monotonic increase over the four years, overall support for the use of the technology in healthcare exhibited a non-linear trajectory. This trajectory rose from initial caution to optimistic acceptance and subsequently recalibrated to a discerning form of critical usage. This recalibration signifies a sophisticated mental model derived from a deeper understanding of the probabilistic nature of the technology and its inherent ethical complexities regarding privacy, accountability, and bias[17].

This work offers three primary contributions to the IUI community. First, it provides a longitudinal empirical account of how user attitudes toward AI technology evolve over an extended period of significant technological and societal change. Second, it proposes a nuanced model of user attitude formation that extends beyond simple acceptance to highlight the pivotal role of critical reflection in shaping a mature human-AI relationship. Finally, it translates these findings into concrete design implications for creating adaptive, transparent, and trustworthy IUIs that support users throughout their evolving usage. To guide our investigation, we pose the following research questions:

- RQ1: What are the evolutionary trajectories of individuals’ attitudes toward the involvement of generative AI-based IUIs in their healthcare?
- RQ2: What key events or factors trigger significant changes in these attitudes?
- RQ3: What are the implications of these evolutionary trajectories for the design of medical IUI systems with long-term adaptability?

2. Related Works

Our research is situated at the intersection of HCI, health informatics, and the emerging field of human-AI interaction. We focus specifically on the longitudinal evolution of user perceptions. This work builds upon and extends three primary streams of literature: the study of trust dynamics in intelligent systems, models of technology acceptance in high-stakes domains, and user-centered research on AI-driven health interfaces. Appendix Table 1 provides an overview of representative categories of AI-based Medical Intelligent User Interfaces. This overview includes their interface types with exemplary

applications, core functionalities, interface characteristics, target user groups, and underpinning AI core technologies.

A central theme in our work is the dynamic nature of trust. This concept has been extensively explored within the HCI and AI communities[18][19]. Early models often treated trust as a relatively static attribute. However, contemporary research emphasizes trust calibration. In this process, the trust of a user in a system appropriately reflects the capabilities, limitations, and context of use of that system[20][21]. This calibration is particularly critical in healthcare. Both under-trust, which leads to the abandonment of useful tools, and over-trust, which leads to the uncritical acceptance of erroneous advice, can have severe consequences[22]. The IUI plays a pivotal role in this process as it serves as the primary channel through which the system communicates competence and uncertainty. Work in Explainable AI (XAI) has shown that interfaces providing insight into the reasoning of the system can improve trust and user performance[23][24][25]. Our findings contribute a crucial longitudinal dimension to this discourse. We demonstrate that trust calibration is not a singular event but a continuous and evolving process. The non-linear pattern of acceptance we observed, where initial optimism yields to critical re-evaluation as user understanding deepens, suggests that effective IUIs must support this continuous process. Systems must move beyond initial user orientation to facilitate ongoing and critical engagement.

Classic frameworks for technology adoption, such as the Technology Acceptance Model (TAM) and its successors, have provided foundational insights by identifying perceived usefulness and ease of use as key drivers of acceptance[26][27]. While valuable, these models often capture a cross-sectional analysis and may not fully account for the complexities of high-stakes domains like healthcare. HCI research has increasingly moved towards nuanced and context-sensitive models that consider the user journey over time[28]. The concept of technology domestication describes how novel technologies are gradually integrated into and transformed by daily routines[29]. Our study empirically illustrates this process in the context of generative AI[30]. The observed shift from users viewing the IUI as a simple information retrieval tool to a sophisticated instrument for preparing for clinical encounters aligns with this perspective. We extend this body of work by showing how increasing AI literacy functions as a key factor in this domestication. This leads to a form of use characterized by discernment and empowerment rather than uncritical adoption.

Finally, a significant body of user-centered research has investigated the specific challenges and opportunities of IUIs in healthcare. Studies on health chatbots[31][32] and diagnostic aids[33] have consistently highlighted key user concerns that align with our findings. These include apprehensions regarding data privacy and real-world consequences such as insurance discrimination, the essential requirement for human oversight and clear accountability structures, and the strong preference for AI to function as an assistant to rather than a replacement for human clinicians[34]. Users often face a fundamental tension. They are drawn to the convenience and accessibility of AI-driven support yet remain wary of potential depersonalization and the loss of empathetic human connection[35]. While existing work has mapped these issues, it has often done so at a single point in time. Our four-year study contributes by revealing how the importance of these concerns evolves. For instance, we found that fears regarding accountability and algorithmic fairness became more pronounced in the later stages of our study. This shift occurred as the understanding of the technology among participants matured and their reliance on it deepened. By charting this evolution, our work provides a detailed understanding of the long-term design requirements for creating responsible and sustainably adopted medical IUIs.

3. Methods

To investigate the dynamic evolution of user attitudes toward generative AI in healthcare (RQ1) and the underlying drivers of these changes (RQ2), we adopted a four-year longitudinal mixed-methods research design spanning from 2022 to 2025. This design enabled the examination of the co-adaptation between users and technology over time. This approach addresses the limitations of static and cross-sectional approaches in capturing interactions with rapidly evolving technologies such as conversa-

tional and generative AI [36]. The mixed-methods approach supported both measurement and interpretation. Validated quantitative scales administered annually tracked changes in AI literacy and attitudes among participants. In-depth semi-structured interviews provided contextual explanations for these shifts. These interviews revealed critical incidents and evolving mental models that underlie the observed attitudinal changes.

3.1. Participants

We initially recruited 21 participants. To capture a diverse range of perspectives beyond the typical tech-savvy university population, we employed a hybrid recruitment strategy. This strategy combined online outreach through social media platforms (WeChat, Red Note, BILIBILI) and offline community visits in Beijing, Tianjin, and Hebei. This approach yielded a sample with varied backgrounds and levels of digital exposure.

The inclusion criteria required participants to be at least 18 years old, possess sufficient cognitive and expressive abilities for the study, and demonstrate prior experience with digital health tools or conversational interfaces. This criterion ensured that participants could meaningfully engage with the evolving technology from the onset of the study in 2022. The exclusion criteria were designed to minimize expert bias and ensure a focus on lay user perspectives. We therefore excluded individuals with cognitive or communication impairments as well as those with direct conflicts of interest, such as professionals or students in AI development or medical informatics.

The study received ethical approval from the Institutional Review Board (IRB) of the university. All individuals were thoroughly briefed on the purpose, procedures, and potential impacts of the study. They were assured of data anonymity and their right to review or delete their data. Participants were also informed of their right to withdraw from the study at any time without penalty. Written informed consent was obtained from all participants before the study commenced. Each participant received a nominal compensation of 30 RMB per year.

The study commenced with an initial cohort of 21 participants. Over the four-year period, the study experienced participant attrition. Two participants withdrew in 2023, two additional participants withdrew in 2024, and four participants withdrew in 2025. This resulted in a final sample of 16 participants who completed the entire longitudinal study. The primary reasons cited for withdrawal were personal scheduling conflicts and loss of contact, rather than factors related to the study topic or trust in the technology. The initial cohort of 21 participants had an age range of 19 to 45 years ($M=27.43$, $SD=7.07$). As detailed in Table 2, the sample represented a diversity of genders, urban and rural backgrounds, educational levels ranging from primary school to bachelor degrees, and professions including students, teachers, editors, drivers, and managers.

3.2. Procedure

Data was collected annually from each participant in 2022, 2023, 2024, and 2025. Each yearly data collection session adhered to a structured two-stage protocol to prioritize quantitative measurement before qualitative exploration. The sequence involved collecting scale data (T1) prior to the interview (T2) to minimize potential priming effects where the reflective discussion might influence self-reported attitude scores. Throughout the longitudinal study, we documented the specific AI tools utilized by participants to track technological adaptation. The reported tools evolved alongside the broader AI landscape, transitioning from early retrieval-based health chatbots and symptom checkers (e.g., Baidu Health, Dingxiang Doctor) in 2022 to advanced generative large language models (e.g., ChatGPT, Ernie Bot) in 2024–2025.

In the first stage (T1), participants independently completed two online questionnaires. These included the IUI and Generative AI Literacy Scale (Appendix Table 3) and the Attitudes toward Generative AI in Healthcare Scale (Appendix Table 4). These instruments quantitatively assessed the self-reported knowledge of participants regarding conversational and generative AI technologies, alongside their evolving perceptions of utility, risk, and trustworthiness.

In the second stage (T2), participants engaged in semi-structured interviews following the protocol outlined in Appendix Table 5 and 6. Interviews were conducted remotely using Tencent Meeting, with durations ranging from 34 to 57 minutes. All sessions were audio-recorded following the acquisition of participant consent. Interviewers completed standardized training in qualitative inquiry and the technical fundamentals of generative AI. They followed the guide while retaining the flexibility to investigate emergent themes. This approach facilitated the exploration of the rationale underlying the quantitative results and captured personal narratives, critical incidents, and evolving mental models regarding medical AI. To ensure data fidelity, audio recordings were transcribed verbatim within 24 hours of each session. Two research team members independently reviewed and proofread each transcript to verify accuracy and consistency.

3.3. Instrument Design

3.3.1. IUI and Generative AI Literacy Scale

To longitudinally track the evolution of participant understanding of conversational and generative AI, we developed the IUI and Generative AI Literacy Scale (Appendix 3). This instrument was designed to assess technical knowledge and systematically evaluate the maturity of the cognitive framework of the user regarding this technology. The mental model of a user regarding how a system works influences interaction strategies, trust, and the overall user experience. The scale comprises six items (A1–A6). Each item targets a distinct facet of this evolving cognitive framework. The instrument employs a 5-point Likert scale ranging from Strongly Disagree to Strongly Agree to capture the intensity of self-assessed literacy.

The items were constructed to represent a progression from foundational awareness to critical understanding. Item A1 establishes a baseline by assessing comprehension of core concepts regarding IUIs and generative AI. Item A2 probes a higher level of understanding by asking participants to differentiate between generative AI and traditional search engines. This distinction reflects an understanding of the shift in interaction paradigms from keyword retrieval to dialogic co-creation. Item A3 measures proactive engagement and continuous learning, which indicate a commitment to updating mental models in response to rapidly changing technology. Item A4 assesses applied practice and gauges whether the user has begun to utilize the technology in meaningful contexts such as learning or research.

The scale also incorporates psychological constructs central to HCI. Item A5 measures self-efficacy and draws from the social cognitive theory of Bandura to assess the confidence of the user in their ability to use the technology effectively. High self-efficacy is a predictor of technology adoption and resilience in the face of errors. Item A6 probes the understanding of system principles, such as the probabilistic nature of the model. This item assesses whether a user has moved beyond an opaque conception of the system to a mental model that acknowledges inherent capabilities and limitations.

To ensure the robustness of this original contribution, we evaluated the reliability and validity of the scale using the collected dataset (N=76). The scale demonstrated excellent internal consistency, with a Cronbach's Alpha of 0.98. The Kaiser-Meyer-Olkin (KMO) measure verified the sampling adequacy for the analysis (KMO = 0.93). Bartlett's test of sphericity indicated that correlations between items were sufficiently large for factor analysis ($p < .001$). A principal component analysis revealed a single-factor structure, with all six items exhibiting high factor loadings ranging from -0.93 to -0.97. These results confirm that the IUI and Generative AI Literacy Scale is a reliable and valid instrument for measuring user cognition in this context.

3.3.2. Attitudes toward Generative AI in Healthcare Scale

To quantitatively measure the evolving attitudes of participants toward the application of generative AI in the medical domain, we developed the Attitudes toward Generative AI in Healthcare Scale (Appendix 4). The design of this instrument aligns with HCI principles and moves beyond classic technology acceptance models to address the complexities of the healthcare domain. Acceptance in this context

involves a trade-off between perceived benefits and perceived risks mediated by context-dependent trust.

The scale contains four interconnected dimensions. Each dimension assesses a specific aspect of the user stance. The first two dimensions are Perceived Usefulness (B1–B4) and Behavioral Acceptance (C1–C4). These extend traditional acceptance frameworks. Perceived Usefulness assesses the degree to which users believe the technology can enhance their health knowledge and management. Behavioral Acceptance measures the willingness to use the IUI in specific scenarios ranging from low-risk tasks such as preparing for a doctor visit to more sensitive applications such as seeking mental health support.

The latter two dimensions are Contextualized Trust (D1–D4) and Perceived Risk (E1–E4). These are central to the HCI approach of the study. The Contextualized Trust dimension facilitates the study of trust calibration rather than treating trust as a single construct. Items D1 and D2 assess foundational trust in the information provided by the AI. Items D3 and D4 specifically probe the understanding of system boundaries and limitations, including the willingness to act without human confirmation. This aligns with the goal of designing systems that foster appropriate trust. The Perceived Risk dimension systematically captures barriers to adoption. These barriers include concerns regarding misinformation, data privacy, and clinical safety. This multi-dimensional design allows for an analysis of the relationship between the perception of utility, trust in outputs, and sensitivity to inherent risks.

We evaluated the psychometric properties of the scale using the valid longitudinal dataset comprising 76 observations. The instrument demonstrated excellent internal consistency across all dimensions. Cronbach's Alpha coefficients were 0.96 for Perceived Usefulness, 0.94 for Behavioral Acceptance, 0.95 for Contextualized Trust, and 0.95 for Perceived Risk. The Kaiser-Meyer-Olkin (KMO) measure was 0.96, verifying sampling adequacy. Bartlett's test of sphericity was significant ($p < .001$), indicating sufficient correlations for factor extraction. An Exploratory Factor Analysis with a four-factor extraction supported the multi-dimensional structure of the instrument. The analysis confirmed that the items effectively differentiate between the theoretical constructs of utility, acceptance, trust, and risk.

3.3.3. Semi-Structured In-depth Interview Protocol for User Attitudes

To complement the quantitative data and examine the contextual factors regarding the evolving attitudes of participants, we designed a semi-structured in-depth interview protocol (Appendix 5 and 6). Within the HCI tradition, qualitative inquiry is essential for moving beyond descriptive perceptions to understanding underlying reasoning. This protocol was developed to identify the nuances, cognitive dissonance, and emotional aspects of the mental model of the user regarding medical AI. These aspects are often inaccessible through structured scales alone.

The protocol follows a progressive structure that guides the participant from general topics to specific and emotionally significant ones. This design intentionally establishes a comfortable context before addressing sensitive topics related to trust and risk. The interview is organized into five modules:

1. Foundation and Context (Q1–Q6): This initial module gathers demographic data and assesses the technological readiness and domain experience of the participant. This provides a contextual basis for interpreting subsequent responses.
2. Priming and Initial Perception (Q7–Q9): This section uses open-ended and associative inquiry regarding concepts such as an AI physician to elicit the initial mental model of the user. It serves to identify baseline assumptions and emotional perspectives before more structured probing begins.
3. Core Exploration (Q10–Q25): This module constitutes the primary focus of the interview and systematically explores five key HCI themes. It investigates Perceived Usefulness by contextualizing the discussion within a complete user journey from symptom onset to recovery (Q10). It examines acceptance boundaries by asking users to define prohibited applications and to weigh trade-offs between convenience and discomfort in sensitive areas such as mental health (Q16). The protocol directly addresses the issue of trust by exploring its formation (Q18), vulnerability

regarding system opacity (Q19), and potential for repair following failure (Q20). Finally, it investigates Risk Perception and Ethics by prompting users to articulate and prioritize concerns regarding privacy (Q23) and to evaluate the complex issue of accountability (Q24).

4. Role Outlook and Emotional Synthesis (Q26–Q28): The concluding module invites participants to synthesize their perspectives using metaphorical descriptions (Q26) and summary statements (Q27). This provides a holistic view of the final stance of the user and captures the emotional dimensions of their relationship with the technology.

This protocol is designed as a flexible framework rather than a rigid script to facilitate in-depth inquiry. The objective is to explain the quantitative findings and generate context-specific insights that can directly inform the design of transparent, trustworthy, and human-centered intelligent user interfaces for healthcare.

3.4. Data Analysis Method

We employed a mixed-methods approach to analyze the data. This allowed us to triangulate findings and develop a comprehensive understanding of the user experience over the four-year period. The quantitative and qualitative data were analyzed in parallel and subsequently integrated. This integration allowed the statistical trends to be explained by the explanatory context derived from the interviews.

For the quantitative data collected from the scales, we calculated the mean score for each of the five core dimensions for every participant in each year. These dimensions included AI Literacy, Perceived Usefulness, Behavioral Acceptance, Contextualized Trust, and Perceived Risk. To analyze the change over time, we conducted a series of one-way repeated measures Analyses of Variance (ANOVAs). This statistical method assesses changes in a continuous dependent variable across multiple time points within the same subject group. The within-subjects factor was Year with four levels corresponding to 2022, 2023, 2024, and 2025. The dependent variables were the mean scores for each of the five dimensions. We applied the Greenhouse-Geisser correction to the degrees of freedom in cases where the assumption of sphericity was violated as indicated by Mauchly's test. For all significant ANOVA results, we performed post-hoc pairwise t-tests with a Bonferroni correction to identify specific years that differed significantly. Statistical significance for all tests was established at an alpha level of $p < .05$.

For the qualitative data obtained from the semi-structured interviews, we conducted a reflexive thematic analysis following the six-phase framework established by Braun and Clarke [37][38]. This method was selected for its flexibility and suitability for identifying patterns of meaning across the dataset. Two researchers analyzed the transcripts through an iterative process. Phase one involved data familiarization through repeated reading of the transcripts from all four years. Phase two involved systematic coding where researchers independently generated initial codes to capture semantic and latent features relevant to the research questions. In phase three, these codes were collated into potential themes. Phases four and five involved a collaborative review process where researchers refined the themes against the dataset and defined clear definitions and names for each theme. This collaborative consensus approach ensured that the final themes were coherent and distinct. The final themes were used to structure the qualitative results and to elucidate the rationale behind the statistical patterns observed in the quantitative analysis.

4. Results

4.1. Quantitative Results

To provide a comprehensive overview of the longitudinal trends, we first present the aggregate trajectories of the four attitude dimensions in Figure 1. The plot illustrates the mean scores and 95% confidence intervals for each year, highlighting the distinct evolutionary paths of user perceptions. Following this

aggregate view, the individual participant trajectories visualized in Figures 2–6 reveal the granular longitudinal patterns across the five core dimensions of the study. The heatmaps demonstrate a consistent linear increase in AI literacy scores. They also illustrate the non-linear trajectories observed in attitudes toward usefulness, acceptance, and trust, alongside the stabilization of risk perceptions over the four-year period.

4.1.1. The Monotonic Rise of AI Literacy

The first hypothesis posited a continuous increase in the AI literacy of participants. The quantitative data supports this hypothesis. The mean AI Literacy score rose steadily from a baseline of ($M = 1.34$, $SD = 0.33$) in 2022 to a high value of ($M = 4.87$, $SD = 0.22$) in 2025. The RM-ANOVA revealed a large and statistically significant effect of Year on AI Literacy scores ($F(3, 45) = 364.45$, $p < .001$, $\eta_g^2 = .92$). Bonferroni-corrected post-hoc tests confirmed that this increase was significant for every consecutive year-to-year comparison including 2022 versus 2023, 2023 versus 2024, and 2024 versus 2025 (all $ps < .001$). This provides statistical evidence for consistent growth in the cognitive understanding and self-efficacy of users regarding generative AI over the four-year period.

4.1.2. The Wave-Like Trajectory of User Attitudes

In contrast to the linear rise in literacy, the three primary attitude dimensions including Perceived Usefulness, Behavioral Acceptance, and Contextualized Trust exhibited non-linear trajectories. The analysis suggests a pattern of increase followed by stabilization or adjustment. RM-ANOVAs showed a significant main effect of Year for all three dimensions: Perceived Usefulness ($F(3, 45) = 160.37$, $p < .001$, $\eta_g^2 = .78$), Behavioral Acceptance ($F(3, 45) = 189.19$, $p < .001$, $\eta_g^2 = .86$), and Contextualized Trust ($F(3, 45) = 166.24$, $p < .001$, $\eta_g^2 = .80$).

Post-hoc tests showed significant increases from 2022 to 2023 and from 2023 to 2024 across all three dimensions (all $ps < .001$). This corresponds to a period of optimistic acceptance. The critical comparison concerned the transition from the 2024 peak to 2025. Perceived Usefulness declined significantly over this interval (from $M = 4.59$, $SD = 0.51$ to $M = 4.34$, $SD = 0.48$; $p = .006$). This provides statistical evidence of a shift in perceived utility. Conversely, Behavioral Acceptance and Contextualized Trust exhibited numerical decreases (from $M = 4.17$ to $M = 3.94$ and from $M = 3.87$ to $M = 3.68$) that did not reach statistical significance.

Overall, the three attitude dimensions demonstrated behaviors distinct from a simple linear increase. While Perceived Usefulness showed a statistically significant decline in the final phase, Behavioral Acceptance and Contextualized Trust indicated a trend toward stabilization. These statistical results provide the structural framework for the qualitative findings, which investigate the contextual reasoning behind these shifts in attitude and the boundaries of trust.

4.1.3. Evolution of Perceived Risk

In the analysis of Perceived Risk, the scoring was inverted to establish a direct correlation between higher values and elevated levels of concern. The RM-ANOVA demonstrated a significant main effect of Year ($F(3, 45) = 133.95$, $p < .001$, $\eta_g^2 = .70$). Post-hoc comparisons indicated that risk concerns peaked in 2022 ($M = 4.77$, $SD = 0.36$). A significant decrease occurred in 2023 ($M = 3.67$, $SD = 0.79$) and subsequently in 2024 ($M = 2.58$, $SD = 0.82$) (all $ps < .001$). However, no significant difference emerged between 2024 and 2025 ($M = 2.68$, $SD = 0.58$). This indicates that following the initial decline, risk perception stabilized. This stability implies that increased user knowledge led to a consistent and defined assessment of technological risks rather than their elimination.

4.2. Qualitative Findings

The thematic analysis of the four-year interview transcripts revealed a comprehensive narrative that elucidates the statistical trends presented above. The evolution of user attitudes toward generative AI

in healthcare is shaped by two distinct yet interconnected trajectories. The first is a consistent longitudinal rise in the AI literacy of participants and the sophistication of their mental models regarding the technology. The second is driven by this cognitive maturation and represents a non-linear evolution of attitudes. This trajectory moves from initial caution to optimistic acceptance and finally to a discerning and critical stance. These themes illustrate the process of technological domestication where users transform from passive recipients into empowered and critical partners in their interaction with medical IUIs.

4.2.1. Progression from Conceptual Ambiguity to Instrumental Utility

Corresponding with the quantitative findings, the qualitative analysis confirmed a continuous and significant increase in the AI literacy of participants from 2022 to 2025. This cognitive growth unfolded across three distinct and successive levels including conceptual understanding, application capability, and critical awareness.

First, the conceptual understanding of participants evolved from vague analogies to a clear and functional mental model. In the early stages of the study, perception was frequently anchored to simpler technologies. As P3 stated in 2022, *"I thought it was just a more advanced Siri, something that could only chat."* By the middle of the study, however, users began to articulate the fundamental paradigm shift represented by generative AI. P11 provided a sharp distinction in 2024: *"Its essential difference from a search engine is creating answers not listing links, which determines that our interaction should be a dialogue, not a keyword battle."* This demonstrates a sophisticated grasp of the interactive and generative capabilities of the IUI.

Second, this conceptual clarity translated into a shift in application capability as users moved from tentative exploration to skilled instrumentalization. Participants evolved from being passive information consumers to proactive problem-solvers who leveraged the IUI as a tool. For instance, P5 described a systematic workflow in 2023: *"I now use it to get preliminary interpretations of obscure medical terms, and even ask it to simulate a doctor's perspective to question the symptom descriptions I've prepared, helping me to identify any gaps."* This instrumentalization of the AI, where it becomes an integrated part of the health management process of a user, established a foundation for initially positive attitudes. By 2025, this proficiency was widespread as 17 participants reported they were very familiar with the process of using generative AI.

Finally, as their understanding deepened, participants developed a conscious and critical awareness of the underlying mechanisms of the technology. They moved beyond an opaque view to a nuanced appreciation of its probabilistic nature. This fostered skepticism and served as a prerequisite for effective trust calibration. The reflection of P6 in 2025 was particularly insightful: *"The more I understand its principle is based on probability and big data, the more clearly I know its boundaries. Its confidence might just be a mimicry, and what I need is for it to also clearly tell me under what circumstances its judgment has the highest uncertainty."* This transition from unconscious acceptance to a critical and informed perspective represents the final stage of cognitive maturation and establishes the context for the subsequent recalibration of their overall attitude.

4.2.2. Non-Linear Evolution of Attitudes: From Optimistic Acceptance to Discerning Recalibration

In contrast to the linear progression of cognition, the support of participants for generative AI in healthcare followed a distinct non-linear trajectory. This pattern was quantitatively confirmed to resemble a progression involving a rise, a peak, and a subsequent recalibration. This suggests that user attitude is not a simple function of knowledge but rather a complex negotiation shaped by different stages of cognitive maturation.

Phase 1 (2022–2024): The Rise to Optimistic Acceptance During the initial two years of the study, the support for and trust in medical IUIs rose significantly as participant understanding grew.

This period was driven primarily by two factors including the demystification of the technology and its successful instrumentalization in daily life. Initially, a lack of understanding generated distrust. However, as users developed a basic mental model of how the AI functioned, this apprehension subsided and was replaced by an appreciation of practical value. As P2 noted in 2024: *"Knowing that it just answers questions by learning from massive amounts of data made it feel less mysterious, and I became more daring to use it."*

This newfound comfort was solidified as participants began to integrate the AI into their health routines and leveraged it as both a tool and a form of digital support. They found tangible value in using the IUI for information consolidation, preparation for clinical visits, and chronic disease management. The AI was framed as a constant and non-judgmental resource. P19 explained in 2024: *"It is like having a personal medical librarian with me at all times who never gets tired. It helps me organize all the information, making me feel more confident when I face a real doctor."*

Phase 2 (2024–2025): The Recalibration to a Discerning Stance This peak of optimism was followed by a statistically significant decrease in support scores between 2024 and 2025. The qualitative data reveals that this was not a reversion to technological resistance but rather the emergence of a mature and critical consciousness triggered once cognitive literacy crossed a specific threshold. As users moved beyond functional utility, they began to engage in deeper ethical and societal reflections.

This critical turn manifested in three key areas. First, there was a reassessment of precision as users began to question the privacy trade-offs inherent in personalization. P18 reflected in 2025: *"The more precise the plan it provides, the more I wonder how much of my data it has collected. And where does that data end up?"* Second, as reliance on the tool deepened, the questioning of responsibility became acute. Abstract utility was weighed against concrete risks. P12 stated in 2025: *"I used to just think it was useful, but now I wonder, what if it is wrong? Who is responsible? If no one takes the blame, then I can not completely entrust myself to it."* Finally, users began a reflection on the human-AI relationship itself. Eleven participants expressed a newfound wariness of over-reliance. P10 articulated this concern in 2025: *"For a while, I asked it almost everything. But then I got scared. I was afraid I would lose the ability and patience to communicate with real doctors, and I also feared the coldness of everything being an algorithm."*

In summary, the decline in support observed in 2025 represents the ultimate expression of user maturation. It marks the transition from naive technological optimism to a sophisticated and critical partnership. The concluding thought of P14 in 2025 encapsulates this nuanced relationship: *"I know it better now than I did four years ago, and I also fear it more. But this fear is not the ignorance and rejection of the past, but a sober respect. I know how much it can do, and more importantly, how much it cannot. I will still use it frequently, but I will no longer trust it without reservation. Perhaps this is the healthiest relationship we can have with AI."*

4.2.3. Domain-Specific Thematic Evolution

Beyond the overarching trajectories, the analysis revealed significant evolutions within specific thematic domains. These shifts regarding how users perceive the usefulness of the IUI, establish boundaries, form trust, and assess risk provide detailed insights for designing future systems.

Transition from Information Retrieval to Proactive Health Management The perception of participants regarding the usefulness of the IUI evolved significantly over the four years. Initially, in 2022, most viewed it as an efficient information retrieval tool. P1 noted: *"I used it to check what medicine to take for a cold. It was faster than Baidu and the explanation was clearer."* This utility was largely confined to the information level. However, by the period between 2023 and 2024, expectations had shifted toward personalized and proactive health management. Users began to desire an IUI that could synthesize data and provide anticipatory insights. P5 described an ideal scenario in 2024: *"I hope that after I input that I have been feeling fatigued lately, it does not just give me a list of possible causes, but combines it with the persistently high resting heart rate data from my wristband to remind me that*

my fatigue might be related to decreased sleep quality and potential stress." This reflects a desire for the IUI to transition from a reactive respondent to a proactive sense-making partner. Despite this, users consistently maintained clear boundaries for the capabilities of the AI and established strict limitations around severe, acute, and emotionally profound health issues. P6 emphasized in 2025: *"I can accept it helping me analyze abnormal indicators in a physical examination report, but if it is persistent chest pain or a suspected tumor, I would never consult an AI. The empathy of a machine is ultimately a program and cannot replace the warm companionship of a human doctor when diagnosing cancer."*

Context-Dependent Acceptance and Human-Centric Boundaries User willingness to engage with the IUI was dependent on the specific context and perceived risk of the task. For low-stakes activities such as light symptom consultation (C1) and preparing for a doctor visit (C4), acceptance was generally high. P11 mentioned in 2023: *"I was very nervous before seeing the dentist, so I used the AI to simulate the questions the doctor might ask. I felt much more prepared."* However, when the context shifted to sensitive domains such as mental health support (C3) or long-term personal data logging (C2), acceptance boundaries became defined. While acknowledging convenience, a distinct discomfort emerged. P17 articulated this tension in 2024: *"Confiding my deepest anxieties and secrets to an algorithm? That sounds pathetic. Behind that convenience is a sense of digital-age loneliness. I might use it for a mindfulness guide, but I would never treat it as an emotional anchor."* This highlights a paradox for IUI design where users desire the efficiency offered by AI but resist applications that they feel might diminish essential human experiences or expose them to clinical risk.

Dynamics of Trust: Transparency, Authority, and Repair Mechanisms The construction of trust was a multi-faceted process. Throughout the study, participants consistently emphasized the importance of transparency and authoritative endorsement as foundational elements of trust. P16 stated in 2025: *"I need to know why it reached this conclusion. If it can show me, step-by-step, which latest clinical guidelines it referenced and its reasoning process, I will trust it more. It would be even better if it were certified by major hospitals like Xiehe or Huashan."* The opacity of many AI systems was identified as a primary barrier to trust. P11 asked in 2023: *"I cannot trust a doctor who cannot even explain the reasons for its judgment. Where is the AI intuition built upon? Data? Which data?"* Findings also revealed that trust is a dynamic and potentially repairable construct. P10 shared an experience in 2024 where trust was restored following a failure: *"Last year, it misidentified the symptoms of a common skin disease, and I was angry. But later, its system proactively pushed a correction notice, explaining the reason for the error and the update. This act made me feel it was alive and responsible, so I was willing to give it another chance."* This suggests that IUIs with transparent error-correction and accountability mechanisms can rebuild user trust.

Privacy Concerns, Accountability, and Algorithmic Fairness Among the various risks discussed, data privacy (E2) was the most prominent concern across all four years. The fears of participants were linked to concrete and negative real-world consequences. P15 explained in 2024: *"What I fear most is not receiving ads, but that my health data will be obtained by insurance companies, leading to me being denied coverage or charged higher premiums in the future. Or worse, my employer finds out and it affects my promotion."* The issue of accountability (Q24) prompted intense discussions. When asked to assign responsibility for an AI-induced medical error, the response of P4 in 2023 represented a desire for a systemic solution: *"It has to be a chain! The developing company must bear the main responsibility, the hospital that introduced it without proper supervision should also be responsible, and the doctor who blindly follows the AI. The patient has the least responsibility because they are in the weakest position."* Furthermore, concerns about algorithmic fairness (Q25) and the potential for AI to exacerbate health disparities grew over time. P5 expressed this concern in 2025: *"Good AI will definitely be given to the upper class first. Technology sometimes widens the gap, it does not narrow it."*

Future Role of the IUI as an Assistive and Augmentative System When asked to describe the ideal role of AI in future healthcare using a metaphor (Q26), participant answers converged on a vision of assistance and augmentation rather than replacement. The metaphors included a tool like a stethoscope (P9, 2023) which emphasized its role as an extension of clinician abilities, a knowledgeable and tireless intern (P14, 2023) which acknowledged information-processing power but implied the need for supervision, and a health navigator (P2, 2025) which highlighted its potential as a guide through a complex healthcare journey. Summarizing the four-year progression, the primary sentiment was one of caution mixed with optimism. The reflection of P8 in 2025 captures this collective perspective: *"I am more and more optimistic about its capabilities. But I am also increasingly vigilant about the problems it might bring. Overall, I think it is an incredibly sharp double-edged sword. What we need to learn is not to throw it away, but how to wield it safely."*

5. Discussion

5.1. Cognitive Maturation and the Emergence of Critical Awareness

In response to RQ1, the longitudinal data reveals a dual trajectory of attitude evolution characterized by the continuous development of instrumental cognition and the subsequent emergence of ethical and societal critique. The positive attitudinal shift from 2022 to 2024 was driven by the progressive development of instrumental cognition, where users transitioned from ambiguous concepts to functional mental models. Participants distinguished the generative nature of the system from traditional search paradigms, leading to increased perceived usefulness and ease of use—core constructs of classic technology acceptance models. Consistent with foundational HCI research on mental models, this functional understanding shaped interaction strategies, allowing users to instrumentalize the AI for efficiency and psychological reassurance while mitigating initial uncertainty [39][40].

However, as cognitive literacy exceeded a critical threshold, critical reflection became the dominant force shaping the 2024–2025 phase, marking a transition from passive reception to empowered, critical partnership. Users shifted focus from functional efficacy to operational mechanisms and societal consequences, aligning with the HCI community's emphasis on value-sensitive design [41]. This critical turn involved examining the technology's fundamental nature: concerns regarding datafication and surveillance emerged alongside apprehensions about accountability in decision support. Furthermore, users renegotiated the human-AI relationship, expressing vigilance against over-reliance and the potential erosion of empathy and moral responsibility inherent in human clinician-patient interactions.

Consequently, the observed decline in support scores in 2025 signifies not technological rejection, but the formation of a mature user mental model—transforming from uncritical optimism to discerning, pragmatic critique. This perspective, which acknowledges AI capabilities while remaining vigilant of risks, constitutes the psychological foundation for sustainable human-AI collaboration in healthcare. Extending theories such as the Computers as Social Actors (CASA) paradigm, these findings demonstrate that deeper cognitive engagement leads to a calibrated partnership rather than unconditional adoption, underscoring the necessity of transparency and ethical design in long-term IUI acceptance [42][43].

5.2. Design Implications for Next-Generation Medical IUIs

In response to RQ3, our longitudinal findings on the evolving user-AI relationship necessitate a fundamental rethinking of design principles for next-generation medical IUIs. The observed trajectory from optimistic acceptance to discerning recalibration indicates that systems must be capable of dynamic adaptation. This adaptation must address user competence as well as evolving ethical sensibilities and trust calibration needs. Future systems must surpass functional intelligence to become contextually adaptive, ethically reflexive, and relationally adept. They should evolve from static tools into collaborative partners that support the longitudinal cognitive and critical development of the user. This necessitates a shift in HCI design paradigms that moves beyond usability and initial acceptance to

foster long-term, sustainable, and ethically grounded human-AI collaboration in healthcare. The following subsections delineate concrete design implications derived from the four-year study and focus on the foundational components of trust calibration, human-AI teaming, embedded ethics, and lifelong co-adaptation.

5.2.1. Designing for Trust Calibration

A paramount design imperative emerging from the data is the need to actively support trust calibration rather than the simplistic maximization of trust. In the high-stakes domain of healthcare, over-trust presents a significant safety hazard that potentially leads to the uncritical adoption of erroneous advice. Conversely, under-trust results in the abandonment of potentially beneficial tools. The progression of participants toward a discerning stance underscores the necessity for IUIs to facilitate an appropriate alignment between user trust and the actual capabilities and limitations of the system. This aligns with the focus of the HCI community on transparent and XAI but extends it by emphasizing the longitudinal and dynamic nature of this process[44]. Design must therefore create interfaces that make the reasoning, confidence, and boundaries of the AI perceptible and understandable throughout the relationship with the technology.

Concretely, this can be operationalized through several HCI-informed strategies. First, explanatory interfaces should move beyond presenting a final answer to visualizing the reasoning process of the system. This includes explicitly communicating uncertainty by presenting differential diagnoses with associated confidence intervals or qualitative indicators such as high, medium, or low certainty. It also involves citing the key evidence or guidelines underpinning conclusions. Furthermore, systems should proactively flag areas where medical knowledge is rapidly evolving or contentious to manage user expectations. Second, visualizing capability boundaries is crucial. The IUI should be designed to explicitly and intuitively communicate its competencies and incompetencies. For example, upon detecting symptom descriptions indicative of acute conditions such as chest pain or severe abdominal pain, the primary response of the system should not be an informational answer but a clear warning and an immediate pathway to human emergency services. Finally, incorporating mechanisms for dynamic trust adjustment can make the system responsive to real-time needs. The IUI could monitor interaction patterns, such as user hesitation or repeated questioning on a topic, and interpret these as signals for required clarity. This could subsequently trigger more detailed explanations or offer an escalation path to a human expert. This approach transforms the IUI from a non-transparent system into a transparent partner. It empowers users to make informed decisions about when and how to rely on outputs, thereby fostering a mature and sustainable form of trust that is essential for long-term adoption in healthcare.

5.2.2. Fostering a Human-AI Teaming Interaction Paradigm

A consistent theme emerging from the longitudinal data is the preference of users for an assistive and augmentative role for AI in healthcare. This view fundamentally opposes a narrative of replacement. The collective user vision articulated through metaphors such as a stethoscope or health navigator necessitates a shift in interaction design paradigms from creating autonomous AI entities to crafting collaborative human-AI teams. This aligns with the HCI discourse on human-AI collaboration and joint activity which emphasizes the importance of shared goals, mutual understanding, and complementary strengths[45][46]. The design challenge is not to maximize the autonomous capabilities of the AI but to optimize the collaborative efficacy of the human-AI dyad. This ensures that the IUI enhances rather than disrupts the agency of the user and pre-existing healthcare relationships[47].

Translating this philosophy into concrete design requires a focus on three key areas. First, explicit role articulation must be embedded into the system architecture. The language, visual design, and interaction flow should consistently reinforce the identity of the system as an assistant. For instance, following an AI-generated analysis of symptoms, the interface could explicitly state that the analysis is intended to support understanding and preparation but that the final diagnosis and treatment plan must be developed in consultation with a healthcare provider. This moves beyond a simple disclaimer and

actively frames the AI output as one component of a larger clinical decision-making process. Second, design should aim to augment and scaffold human-human interaction, particularly the clinician-patient relationship. Instead of bypassing the doctor, the IUI can be designed as a tool that empowers the patient for more productive clinical encounters. For example, it could help users generate a structured summary of their symptom history, a timeline of key events, or a list of prepared questions for their appointment. This effectively makes the patient a better-informed partner in the consultation. This approach resonates with the HCI concept of interactional empowerment where technology equips users with better resources for human communication. Finally, providing seamless and low-friction transition mechanisms to human experts is a critical feature for maintaining the appropriate boundary and user safety. When a user query exceeds the competence of the system or when the user expresses a desire for human contact, the IUI should facilitate the transition. This could involve an immediate option to connect to a human-operated chat service, schedule a telehealth call, or be provided with direct contact information for relevant support lines. By designing for this collaborative continuum, we create IUIs that users are more likely to trust and integrate sustainably as they respect and reinforce the value of human connection and expertise in care.

5.2.3. Embedding Ethics and Privacy as Foundational Infrastructure

The persistent anxieties surrounding data privacy, accountability, and algorithmic fairness uncovered in the data necessitate a paradigmatic shift in how these concerns are addressed in medical IUI design. Rather than treating ethics and privacy as peripheral compliance issues, they must be elevated to the status of core architectural components. This fundamentally shapes the interaction logic and data flows of the system from the foundational level. This aligns with the HCI framework of Value-Sensitive Design (VSD) and the principle of Privacy by Design which advocate for proactively embedding moral and ethical values into the technical specification[48][49]. In the context of healthcare IUIs, the system architecture must be intrinsically oriented towards preserving user autonomy, ensuring justice, and mitigating harm.

This foundational integration can be operationalized through several key design strategies. First, regarding privacy and data governance, the IUI must transcend simple consent dialogs and embrace comprehensive transparency and user control. This involves implementing the principle of data minimization where only strictly necessary data is collected for a specific and user-understood purpose. Furthermore, the interface should provide an intuitive and visually clear data lineage dashboard allowing users to see what data has been collected, how it is being used, and with whom it has been shared. Critically, this must be coupled with granular data rights including the ability to view, export, and permanently delete personal data with minimal effort. This approach directly addresses concrete fears regarding insurance and employment discrimination and transforms the IUI from a data extractor into a data steward. Second, the issue of accountability must be made interactionally explicit. Prior to users engaging with high-stakes functionalities, the system should employ interactive protocols that clearly delineate the boundaries of responsibility among the various actors including the AI developer, the healthcare institution, the clinician, and the user. This fosters a shared mental model of the socio-technical system. Finally, to combat algorithmic bias and ensure fairness, the IUI architecture must incorporate bias detection and mitigation protocols. Moreover, the interaction design itself must be inherently inclusive by offering multi-modal interfaces to accommodate diverse user populations including the elderly and those with lower health or digital literacy. This ensures that the benefits of the technology are distributed equitably.

5.2.4. Lifelong Learning Systems for Co-Adaptation

The longitudinal findings demonstrate that the relationship between a user and a medical IUI is a dynamically evolving process characterized by significant growth in cognitive understanding and critical awareness. This necessitates a fundamental shift from designing static interfaces to developing lifelong learning systems capable of co-adapting with the user concurrently. This concept aligns with the HCI

vision of co-adaptive systems where the interface and the user mutually influence and reshape each other through sustained interaction[50][51][52]. A medical IUI must therefore possess the capacity to sense, interpret, and respond to the evolving mental model, literacy level, and informational needs of the user.

Operationalizing this co-adaptation requires the implementation of several key interactive mechanisms. First, the system should feature personalized information depth by dynamically adjusting the complexity and nature of its explanations and outputs in response to demonstrated literacy. For a novice user, explanations might focus on actionable advice and simple language. As the system detects increasing sophistication through interaction patterns, it can progressively introduce more detailed technical rationales, discuss underlying uncertainties, and contextualize information within broader medical debates. Second, dynamic risk communication is essential. The manner in which potential risks and limitations are communicated should evolve alongside the user. Initial interactions might present concise warnings while experienced user scenarios could trigger a discussion that includes the probabilistic nature of the AI suggestion and an elaboration of the ethical trade-offs involved. Finally, establishing a participatory design feedback loop directly within the IUI ecosystem transforms the user from a passive consumer into an active co-designer. The system should provide channels for users to report errors and contribute critical reflections on interface elements, explanation styles, or ethical concerns. These insights must then be systematically aggregated and fed back into the iterative development cycle. By designing for this continuous and bidirectional adaptation, medical IUIs can become learning partners capable of growing in sophistication and value alongside the users they serve.

5.3. Theoretical Contributions

The findings of this longitudinal study necessitate a significant extension of classical technology acceptance models, such as the Technology Acceptance Model (TAM), which rely on static constructs of usefulness and ease of use. These frameworks prove insufficient to capture the dynamic, non-linear, and critically reflexive nature of the user-IUI relationship in high-stakes healthcare contexts[53][54]. Consequently, we propose a Critical-Dynamic Framework of Technology Acceptance for medical IUIs. This framework views acceptance not as a static outcome derived from functional utility, but as a continuous process of calibration and negotiation. It integrates key HCI concepts to explain the complex interaction between cognitive evolution and attitudinal shifts, focusing on the sustainability and ethical quality of the human-AI partnership.

Aligned with Value-Sensitive Design (VSD) and co-adaptive systems, our framework provides empirical evidence of how user values—such as privacy and accountability—become increasingly salient through cognitive maturation, thereby driving co-adaptation[55][56]. Furthermore, it refines the CASA paradigm by demonstrating that deep, sustained interaction in high-stakes domains does not lead to unconditional social acceptance, but rather to a critical and calibrated partnership.

The framework introduces three pivotal constructs essential for long-term acceptance. First, Critical Literacy is posited as a time-variant variable extending beyond functional knowledge to include socio-technical understanding. This literacy exhibits a dual-phasic impact: initially driving adoption through demystification, but eventually triggering attitudinal recalibration as critical awareness deepens, aligning with research on mental models and algorithmic literacy[57][58][59]. Second, Contextualized Trust is reconceptualized as a dynamic, multi-dimensional construct requiring continuous negotiation through both functional experiences and value-driven reflections[60][61].

Finally, Perceived Risk expands from a simple barrier into an evolving cluster of concerns, encompassing informational, clinical, and ethical dimensions such as relational erosion and algorithmic bias. The salience of these risks evolves alongside critical literacy, becoming more concrete over time[62][63][64]. Synthesizing these constructs, our framework frames technology acceptance as a longitudinal journey of dynamic trade-offs between instrumental utility and profound human values[65][66]. This perspective challenges linear progression models, emphasizing that sustainable acceptance is rooted not in uncritical reliance, but in a mature, critically informed, and dynamically managed partnership.

5.4. Limitations and Future Work

First, the cultural and geographical specificity of our participant sample, drawn from specific regions in China, constrains the generalizability of the proposed Critical-Dynamic Framework [67]. As cultural factors profoundly shape technology adoption and trust dynamics, future HCI research should pursue cross-cultural longitudinal studies to examine how the observed non-linear attitudinal trajectories and trust calibration mechanisms manifest across diverse societal contexts [68][69]. Such comparative work is essential for developing culturally attuned design principles for global health IUI deployment.

Second, while a four-year period is substantial, it may remain insufficient to fully capture the long-term societal integration and normalization of disruptive generative AI technologies. The observed recalibration phase might represent a transitional fluctuation rather than a permanent state. Future inquiries should aim for longer-term studies to track how attitudes stabilize, transform, or are disrupted by technological and regulatory shifts [70]. This approach aligns with the HCI community's focus on the temporal unfolding of user experiences beyond initial adoption.

Third, the study's singular focus on lay users highlights the need for a broader investigation into the human-AI ecosystem involving clinicians, administrators, and policymakers. Future research should adopt multi-stakeholder and participatory design approaches to investigate the complex interplay between these actors, particularly how patient IUI usage impacts clinician workflows and the doctor-patient relationship. Exploring these dynamics through Computer-Supported Cooperative Work (CSCW) frameworks could yield crucial insights for designing systems that support collaborative care models rather than creating informational asymmetries.

Finally, we acknowledge that the repeated in-depth interview process may have functioned as an intervention, inducing reflection that influenced participant engagement. Future longitudinal studies should explicitly account for such methodological influences through reflexive journals or pre- and post-study assessments. Furthermore, methodological innovation is needed to complement mixed-methods approaches; integrating robust computational interaction logs with qualitative data and conducting experimental studies on specific design interventions—such as uncertainty visualization—would provide a more granular and actionable understanding of trust calibration across different literacy levels.

6. Conclusion

This four-year longitudinal study establishes a dual-trajectory model of user engagement with medical IUIs, where monotonic growth in AI literacy contrasts with a non-linear engagement path evolving from optimistic adoption to discerning recalibration. Challenging linear technology acceptance models, these findings indicate that sustainable high-stakes interaction relies on a mature, critical partnership where utility is continuously weighed against ethical risks. Consequently, medical IUI design must shift to prioritize trust calibration over maximization, reinforce human-AI teaming through explicit role articulation, and embed ethics as core architectural elements. Contributing the Critical-Dynamic Framework to HCI theory, this work recontextualizes acceptance as an ongoing negotiation between instrumental utility and human values, framing medical AI integration as a profound design challenge aimed at cultivating critically engaged user empowerment.

7. Declaration on Generative AI

During the preparation of this manuscript, we utilized large language models in a limited capacity solely for editorial refinement and consistency checks. All authors assume full responsibility for the content, analysis, and interpretations presented herein. All data are authentic, and no AI-generated or artificially fabricated material has been incorporated into the substantive portions of this work.

References

- [1] N. N. Ontika, H. A. Syed, S. M. Saßmannshausen, R. H. Harper, Y. Chen, S. Y. Park, M. Grisot, A. Chow, N. Blaumer, A. F. Pinatti de Carvalho, et al., Exploring human-centered ai in healthcare: diagnosis, explainability, and trust, in: *Proceedings of 20th European Conference on Computer-Supported Cooperative Work, European Society for Socially Embedded Technologies (EUSSET)*, 2022.
- [2] T. O. Andersen, F. Nunes, L. Wilcox, E. Coiera, Y. Rogers, Introduction to the special issue on human-centred ai in healthcare: Challenges appearing in the wild, 2023.
- [3] S. R. Sindiramutty, K. R. V. Prabakaran, R. Akbar, M. Hussain, N. A. Malik, Generative ai for secure user interface (ui) design, in: *Reshaping CyberSecurity With Generative AI Techniques*, IGI Global Scientific Publishing, 2025, pp. 333–394.
- [4] P. Brusilovsky, D. Parra, B. Rahdari, S. Raj, H. Torkamaan, Healthiui: Workshop on intelligent and interactive health user interfaces, in: *Companion Proceedings of the 30th International Conference on Intelligent User Interfaces*, 2025, pp. 183–186.
- [5] S. Dai, J. Davidson, B. Ullmer, W. E. Newman, M. K. Konkel, Generative ai syntheses of platform, content, visuals, and kinetics for cyberphysical computationally-mediated posters and broader applications, in: *Companion Proceedings of the 29th International Conference on Intelligent User Interfaces*, 2024, pp. 45–49.
- [6] J. Yeon, Y. Jung, Y. Baek, D. Lee, J. Shin, W. Y. Chung, User preferences on a generative ai user interface through a choice experiment, *International Journal of Human–Computer Interaction* 41 (2025) 7626–7637.
- [7] Y. Cao, P. Jiang, H. Xia, Generative and malleable user interfaces with generative and evolving task-driven data model, in: *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, 2025, pp. 1–20.
- [8] S. T. Völkel, C. Schneegass, M. Eiband, D. Buschek, What is” intelligent” in intelligent user interfaces? a meta-analysis of 25 years of iui, in: *Proceedings of the 25th international conference on intelligent user interfaces*, 2020, pp. 477–487.
- [9] M. Bhat, Designing ai interfaces for transparent decision-making and ethical reflection, in: *Companion Proceedings of the 30th International Conference on Intelligent User Interfaces*, 2025, pp. 211–214.
- [10] X. Sun, Y. Liu, J. A. Bosch, Z. Li, Interface matters: Exploring human trust in health information from large language models via text, speech, and embodiment, *Proceedings of the ACM on Human-Computer Interaction* 9 (2025) 1–32.
- [11] J. Hermann, N. Jansen, A. Dogangün, We need to understand where the data comes from: Co-designing transparent ai systems for caregiving, in: *Companion Proceedings of the 30th International Conference on Intelligent User Interfaces*, 2025, pp. 105–109.
- [12] S. Mehrotra, C. Degachi, O. Vereschak, C. M. Jonker, M. L. Tielman, A systematic review on fostering appropriate trust in human-ai interaction: Trends, opportunities and challenges, *ACM Journal on Responsible Computing* 1 (2024) 1–45.
- [13] R. A. Leist, H.-J. Profitlich, T. Hunsicker, D. Sonntag, Towards trustable intelligent clinical decision support systems: A user study with ophthalmologists, in: *Proceedings of the 30th International Conference on Intelligent User Interfaces*, 2025, pp. 1470–1484.
- [14] W. Vossen, M. Szymanski, K. Verbert, The effect of personalizing a psychotherapy conversational agent on therapeutic bond and usage intentions, in: *Proceedings of the 29th International Conference on Intelligent User Interfaces*, 2024, pp. 761–771.
- [15] P. Rajpurkar, E. Chen, O. Banerjee, E. J. Topol, Ai in health and medicine, *Nature medicine* 28 (2022) 31–38.
- [16] J. Qiu, L. Li, J. Sun, J. Peng, P. Shi, R. Zhang, Y. Dong, K. Lam, F. P.-W. Lo, B. Xiao, et al., Large ai models in health informatics: Applications, challenges, and the future, *IEEE Journal of Biomedical and Health Informatics* 27 (2023) 6074–6087.
- [17] I. Dirgová Luptáková, J. Pospíchal, L. Huraj, Beyond code and algorithms: Navigating ethical

complexities in artificial intelligence, in: *Proceedings of the Computational Methods in Systems and Software*, Springer, 2023, pp. 316–332.

- [18] E. LaRosa, D. Danks, Impacts on trust of healthcare ai, in: *Proceedings of the 2018 AAAI/ACM Conference on AI, Ethics, and Society*, 2018, pp. 210–215.
- [19] C. Lee, K. Cha, Toward the dynamic relationship between ai transparency and trust in ai: a case study on chatgpt, *International Journal of Human–Computer Interaction* 41 (2025) 8086–8103.
- [20] R. Rosenbacke, Å. Melhus, M. McKee, D. Stuckler, How explainable artificial intelligence can increase or decrease clinicians’ trust in ai applications in health care: systematic review, *Jmir Ai* 3 (2024) e53207.
- [21] J. Kauttonen, R. Rousi, A. Alamäki, Trust and acceptance challenges in the adoption of ai applications in health care: Quantitative survey analysis, *Journal of Medical Internet Research* 27 (2025) e65567.
- [22] K. Kostick-Quenet, B. H. Lang, J. Smith, M. Hurley, J. Blumenthal-Barby, Trust criteria for artificial intelligence in health: normative and epistemic considerations, *Journal of medical ethics* 50 (2024) 544–551.
- [23] Y. Rong, T. Leemann, T.-T. Nguyen, L. Fiedler, P. Qian, V. Unhelkar, T. Seidel, G. Kasneci, E. Kasneci, Towards human-centered explainable ai: A survey of user studies for model explanations, *IEEE transactions on pattern analysis and machine intelligence* 46 (2023) 2104–2122.
- [24] S. Mohseni, N. Zarei, E. D. Ragan, A multidisciplinary survey and framework for design and evaluation of explainable ai systems, *ACM Transactions on Interactive Intelligent Systems (TiiS)* 11 (2021) 1–45.
- [25] P. Wang, H. Ding, The rationality of explanation or human capacity? understanding the impact of explainable artificial intelligence on human-ai trust and decision performance, *Information Processing & Management* 61 (2024) 103732.
- [26] F.-Y. Pai, K.-I. Huang, Applying the technology acceptance model to the introduction of healthcare information systems, *Technological forecasting and social change* 78 (2011) 650–660.
- [27] A. S. Ahadzadeh, S. P. Sharif, F. S. Ong, K. W. Khong, Integrating health belief model and technology acceptance model: an investigation of health-related internet use, *Journal of medical Internet research* 17 (2015) e3564.
- [28] N. Sambasivan, S. Kapania, H. Highfill, D. Akrong, P. Paritosh, L. M. Aroyo, “everyone wants to do the model work, not the data work”: Data cascades in high-stakes ai, in: *proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*, 2021, pp. 1–15.
- [29] R. Williams, S. Anderson, K. Cresswell, M. S. Kannelønning, H. Mozaffar, X. Yang, Domesticating ai in medical diagnosis, *Technology in Society* 76 (2024) 102469.
- [30] S. Breuer, M. Braun, D. Tigard, A. Buyx, R. Müller, How engineers’ imaginaries of healthcare shape design and user engagement: A case study of a robotics initiative for geriatric healthcare ai applications, *ACM Transactions on Computer-Human Interaction* 30 (2023) 1–33.
- [31] A. A. Abd-Alrazaq, M. Alajlani, N. Ali, K. Denecke, B. M. Bewick, M. Househ, Perceptions and opinions of patients about mental health chatbots: scoping review, *Journal of medical Internet research* 23 (2021) e17828.
- [32] A. Yuan, E. Garcia Colato, B. Pescosolido, H. Song, S. Samtani, Improving workplace well-being in modern organizations: A review of large language model-based mental health chatbots, *ACM Transactions on Management Information Systems* 16 (2025) 1–26.
- [33] J. L. Müskens, R. B. Kool, S. A. van Dulmen, G. P. Westert, Overuse of diagnostic testing in healthcare: a systematic review, *BMJ quality & safety* 31 (2022) 54–63.
- [34] M. K. K. Rony, M. R. Parvin, M. Wahiduzzaman, M. Debnath, S. D. Bala, I. Kayesh, “i wonder if my years of training and expertise will be devalued by machines”: concerns about the replacement of medical professionals by artificial intelligence, *SAGE open nursing* 10 (2024) 23779608241245220.
- [35] Z. Zhang, J. Wang, Can ai replace psychotherapists? exploring the future of mental health care, *Frontiers in psychiatry* 15 (2024) 1444382.
- [36] M. J. Brandt, G. S. Morgan, Between-person methods provide limited insight about within-person belief systems., *Journal of Personality and Social Psychology* 123 (2022) 621.

- [37] V. Braun, V. Clarke, Using thematic analysis in psychology, *Qualitative research in psychology* 3 (2006) 77–101.
- [38] V. Braun, V. Clarke, *Thematic analysis: A practical guide* (2021).
- [39] J. H. Gerlach, F.-Y. Kuo, Understanding human-computer interaction for information systems design, *MIS quarterly* (1991) 527–549.
- [40] S. Benford, G. Giannachi, B. Koleva, T. Rodden, From interaction to trajectories: designing coherent journeys through user experiences, in: *Proceedings of the SIGCHI conference on human factors in computing systems*, 2009, pp. 709–718.
- [41] M. Sadek, C. Mougenot, Challenges in value-sensitive ai design: Insights from ai practitioner interviews, *International Journal of Human–Computer Interaction* 41 (2025) 10877–10894.
- [42] J. Ham, S. Li, J. Looi, M. S. Eastin, Virtual humans as social actors: Investigating user perceptions of virtual humans’ emotional expression on social media, *Computers in Human Behavior* 155 (2024) 108161.
- [43] K. Xu, X. Chen, L. Huang, Deep mind in social responses to technologies: A new approach to explaining the computers are social actors phenomena, *Computers in Human Behavior* 134 (2022) 107321.
- [44] U. Ehsan, P. Wintersberger, Q. V. Liao, E. A. Watkins, C. Manger, H. Daumé III, A. Riener, M. O. Riedl, Human-centered explainable ai (hcxai): beyond opening the black-box of ai, in: *CHI conference on human factors in computing systems extended abstracts*, 2022, pp. 1–7.
- [45] P. Hemmer, M. Schemmer, N. Kühl, M. Vössing, G. Satzger, Complementarity in human-ai collaboration: Concept, sources, and evidence, *European Journal of Information Systems* (2025) 1–24.
- [46] N. Jiang, X. Liu, H. Liu, E. T. K. Lim, C.-W. Tan, J. Gu, Beyond ai-powered context-aware services: the role of human–ai collaboration, *Industrial Management & Data Systems* 123 (2023) 2771–2802.
- [47] S. Pareek, N. van Berkel, E. Velloso, J. Goncalves, Effect of explanation conceptualisations on trust in ai-assisted credibility assessment, *Proceedings of the ACM on Human-Computer Interaction* 8 (2024) 1–31.
- [48] R. Y. Wong, D. K. Mulligan, Bringing design to the privacy table: Broadening “design” in “privacy by design” through the lens of hci, in: *Proceedings of the 2019 CHI conference on human factors in computing systems*, 2019, pp. 1–17.
- [49] R. Y. Wong, D. K. Mulligan, E. Van Wyk, J. Pierce, J. Chuang, Eliciting values reflections by engaging privacy futures using design workbooks, *Proceedings of the ACM on Human-Computer Interaction* 1 (2017) 1–26.
- [50] M. Beaudouin-Lafon, S. Bødker, W. E. Mackay, Generative theories of interaction, *ACM Transactions on Computer-Human Interaction (TOCHI)* 28 (2021) 1–54.
- [51] P. Gallina, N. Bellotto, M. Di Luca, Progressive co-adaptation in human-machine interaction, in: *2015 12th International Conference on Informatics in Control, Automation and Robotics (ICINCO)*, volume 2, IEEE, 2015, pp. 362–368.
- [52] M. Kulkarni, M. P. Angurula, Trends in computer science and information technology research, *International Journal of Human–Computer Interaction* 36 (2024).
- [53] A. K. Yarbrough, T. B. Smith, Technology acceptance among physicians: a new take on tam, *Medical care research and review* 64 (2007) 650–672.
- [54] J. D. Portz, E. A. Bayliss, S. Bull, R. S. Boxer, D. B. Bekelman, K. Gleason, S. Czaja, Using the technology acceptance model to explore user experience, intent to use, and use behavior of a patient portal among older adults with multiple chronic conditions: descriptive qualitative study, *Journal of medical Internet research* 21 (2019) e11604.
- [55] S. Ahlberg, A. Axelsson, P. Yu, W. S. Cortez, Y. Gao, A. Ghadirzadeh, G. Castellano, D. Kragic, G. Skantze, D. V. Dimarogonas, Co-adaptive human–robot cooperation: summary and challenges, *Unmanned Systems* 10 (2022) 187–203.
- [56] B. Friedman, P. H. Kahn Jr, A. Borning, A. Hultgren, Value sensitive design and information systems, in: *Early engagement and new technologies: Opening up the laboratory*, Springer, 2013, pp. 55–95.
- [57] D. Shin, A. Rasul, A. Fotiadis, Why am i seeing this? deconstructing algorithm literacy through

- the lens of users, *Internet Research* 32 (2022) 1214–1234.
- [58] A. Oeldorf-Hirsch, G. Neubaum, What do we know about algorithmic literacy? the status quo and a research agenda for a growing field, *New Media & Society* 27 (2025) 681–701.
 - [59] D. E. Silva, C. Chen, Y. Zhu, Facets of algorithmic literacy: Information, experience, and individual factors predict attitudes toward algorithmic systems, *New Media & Society* 26 (2024) 2992–3017.
 - [60] M. Naiseh, D. Al-Thani, N. Jiang, R. Ali, How the different explanation classes impact trust calibration: The case of clinical decision support systems, *International Journal of Human-Computer Studies* 169 (2023) 102941.
 - [61] M. Naiseh, D. Al-Thani, N. Jiang, R. Ali, Explainable recommendation: when design meets trust calibration, *World Wide Web* 24 (2021) 1857–1884.
 - [62] K. Patel, Ethical reflections on data-centric ai: balancing benefits and risks, Available at SSRN 4993089 (2024).
 - [63] M. Corfmat, J. T. Martineau, C. Régis, High-reward, high-risk technologies? an ethical and legal account of ai development in healthcare, *BMC medical ethics* 26 (2025) 4.
 - [64] E. Ratti, M. Morrison, I. Jakab, Ethical and social considerations of applying artificial intelligence in healthcare—a two-pronged scoping review, *BMC Medical Ethics* 26 (2025) 68.
 - [65] K. S. Gill, Prediction paradigm: the human price of instrumentalism, *AI & society* 35 (2020) 509–517.
 - [66] A. Jedličková, Ethical approaches in designing autonomous and intelligent systems: a comprehensive survey towards responsible development, *AI & society* 40 (2025) 2703–2716.
 - [67] D. Kember, L. Gow, Cultural specificity of approaches to study, *British Journal of Educational Psychology* 60 (1990) 356–363.
 - [68] P. A. Atroszko, C. S. Andreassen, M. D. Griffiths, S. Pallesen, Study addiction: A cross-cultural longitudinal study examining temporal stability and predictors of its changes, *Journal of behavioral addictions* 5 (2016) 357–362.
 - [69] C. Lu, S. Wan, Z. Liu, Determinants of depressive symptoms in multinational middle-aged and older adults, *NPJ Digital Medicine* 8 (2025) 501.
 - [70] A. Terracciano, Secular trends and personality: Perspectives from longitudinal and cross-cultural studies—commentary on trzesniewski & donnellan (2010), *Perspectives on Psychological Science* 5 (2010) 93–96.

Appendix

Table 1
Overview of AI-based Medical Intelligent User Interfaces.

Interface Type / Representative Applications	Core Functions & Interface Characteristics	Target Users	Core AI Technologies
AI Health Chatbots (e.g., Ada Health, Babylon Health, K Health)	1. Intelligent Symptom Triage: Simulates a physician's consultation logic through dynamic, multi-turn dialogue (e.g., inquiring about symptom details, duration, aggravating/alleviating factors) to progressively narrow down potential issues. 2. Personalized Risk Assessment: Integrates user-provided health profiles, age, gender, and other data to offer personalized risk analysis rather than generic lists. 3. Triage & Navigation Recommendations: Provides clear next steps, such as "monitor at home," "book a GP appointment within 3 days," or "seek immediate emergency care," with direct links to nearby medical facilities. 4. Health Knowledge Dissemination: Offers authoritative and easy-to-understand explanations of diseases, medications, and healthy lifestyles in a Q&A format.	Patients, general users, and individuals seeking preliminary medical information.	Natural Language Processing (NLU/NLG), Large-Scale Knowledge Graphs, Bayesian Inference Algorithms.
AI Mental Health Companions (e.g., Woebot, Wysa, Youper)	1. Mood Tracking & Insights: Guides users to log daily emotions and visualizes trends through charts to help identify emotional triggers. 2. Structured Intervention Tools: Delivers tools based on CBT (Cognitive Behavioral Therapy) and ACT (Acceptance and Commitment Therapy), such as mindfulness exercises, thought records, and guided breathing, with instructions for use. 3. Empathetic Dialogue: Utilizes techniques like active listening, empathetic feedback, and open-ended questions to create a safe space for expression, rather than offering simple reassurances. 4. Crisis Detection & Intervention: Automatically triggers a crisis response protocol upon detecting expressions of self-harm or suicide, providing critical resources like emergency hotlines.	Individuals experiencing emotional distress or high stress, as well as general users seeking to improve their psychological resilience.	Affective Computing, Sentiment Analysis, Cognitive Behavioral Therapy Models, Dialogue State Management.
Personalized Health Assistants (e.g., Google/Apple Health platforms, Fitbit App)	1. Multi-Source Data Fusion: Seamlessly integrates fragmented data from smartwatches, blood pressure monitors, sleep trackers, manual inputs, and Electronic Health Records (EHR) into a unified health dashboard. 2. Intelligent Trend Detection & Alerts: Proactively identifies abnormal patterns in health data (e.g., sustained increase in resting heart rate, significant decline in sleep quality) and provides notifications with potential explanations. 3. Actionable Insights & Goal Management: Translates data into executable advice (e.g., "Based on your activity level, your recommended daily water intake is 2.5 liters") and helps users set and track personal health goals (e.g., weight loss, step count). 4. Research Participation & Contribution: Allows users to anonymously donate their health data for medical research, contributing to public health studies and disease research (e.g., Apple Heart Study).	Fitness enthusiasts, health-conscious individuals, and patients with chronic conditions (for daily monitoring).	Machine Learning, Time-Series Analysis, Pattern Recognition, Data Visualization.

Table 2

Demographic Characteristics of Participants (2022-2025). Note: ID denotes the participant identification number. "A" represents Age, and "G" indicates Gender. "E" represents Education level, with possible values: B (Bachelor's), H (High School), J (Junior High School), and P (Primary School). "R" (2022, 2023, 2024, 2025) denotes Residence, with two possible values: R (Rural) and U (Urban). "O" stands for Occupation. The symbol "*" indicates absence or missing data.

ID	G	A	2022			2023			2024			2025		
			R	O	E	R	O	E	R	O	E	R	O	E
1	F	24	U	Coach	B	U	Coach	B	U	Coach	B	U	Coach	B
2	F	25	U	Author	B	U	Author	B	U	Author	B	U	Editor	B
3	F	25	U	Cleaning	B	R	Cleaning	B	R	Cleaning	B	U	Cleaning	B
4	M	25	U	Cleaning	J	U	Cleaning	J	U	Cleaning	J	U	Delivery	J
5	M	29	R	Editor	B	R	Editor	B	R	Editor	B	R	Editor	B
6	F	31	R	Professor	B	R	Professor	B	R	Professor	B	R	Professor	B
7	M	33	U	Editor	B	U	Editor	B	R	Teacher	B	*	*	*
8	F	19	R	Student	H	R	Student	H	R	Student	H	R	Teacher	B
9	M	25	R	Coach	B	R	Coach	B	R	Coach	B	R	Coach	B
10	M	27	R	Teacher	B	R	Teacher	B	R	Teacher	B	R	Teacher	B
11	M	22	R	Student	H	R	Editor	B	R	Editor	B	R	Reporter	B
12	F	24	U	Chef	B	U	Chef	B	U	Chef	B	U	Chef	B
13	F	33	R	Driver	P	R	Driver	P	R	Driver	P	R	Cleaning	B
14	M	34	R	Editor	B	R	Editor	B	R	Editor	B	R	Editor	B
15	F	45	U	Labor	H	*	*	*	U	Labor	H	*	*	*
16	F	43	U	Manager	H	U	Manager	High	U	Manager	H	U	Manager	H
17	M	26	R	Pilot	B	R	Pilot	B	U	Pilot	B	R	Designer	B
18	M	25	U	Secretary	B	*	*	*	*	*	*	U	Secretary	B
19	M	19	U	Student	H	U	Student	H	U	Student	B	U	Student	B
20	F	18	R	Student	H	R	Student	H	R	Student	H	*	*	*
21	F	24	R	Translator	B	R	Translator	B	*	*	*	*	*	*

Table 3

IUI and Generative AI Literacy Scale.

Item	Statement	Strongly Disagree	Disagree	Neutral	Agree	Strongly Agree
A1	I understand the basic concepts of IUI (particularly generative AI-driven IUIs).	1	2	3	4	5
A2	I can clearly differentiate between generative AI-driven IUIs and traditional search engines (e.g., Google).	1	2	3	4	5
A3	I actively follow news and discussions concerning the development of generative AI.	1	2	3	4	5
A4	I have utilized generative AI for academic or research purposes.	1	2	3	4	5
A5	I consider myself to be a proficient user of generative AI.	1	2	3	4	5
A6	I believe I have a sufficient understanding of the working principles of generative AI.	1	2	3	4	5

Table 4

Attitudes toward Generative AI in Healthcare Scale. All items are rated on a 5-point Likert scale from Strongly Disagree (1) to Strongly Agree (5).

Dimension	Item	Statement
Perceived Usefulness	B1	I believe using generative AI can help me quickly understand basic medical and health knowledge.
	B2	I believe AI can help me better understand complex medical concepts explained by a doctor.
	B3	I find AI-generated personalized health plans (e.g., diet, exercise) to be instructive.
	B4	I believe using AI for daily health monitoring or consultation would make my life more convenient.
Behavioral Acceptance	C1	For minor, non-urgent health issues, I am willing to first seek advice from an AI.
	C2	I am willing to use an AI to help me record and analyze my long-term health data (e.g., sleep, heart rate).
	C3	When I feel down or stressed, I am willing to try conversing with an AI for psychological support.
	C4	Before a doctor's appointment, I am willing to use an AI to help me organize and prepare the symptoms I need to describe.
Contextualized Trust	D1	Overall, I believe the medical and health information provided by generative AI is reliable.
	D2	I believe AI is accurate in analyzing medical data (such as symptom descriptions).
	D3	If an AI provides a recommendation, I would feel comfortable acting on it without seeking confirmation from a human doctor.
	D4	I believe an AI can understand my unique emotional and life context as a person.
Perceived Risk	E1	I am not concerned that an AI might provide me with incorrect or outdated medical information.
	E2	I am not concerned that my personal privacy and health data will be compromised when using AI health services.
	E3	I am not concerned that over-reliance on an AI would cause me to ignore my body's real signals or delay seeking medical attention.
	E4	Discussing my health issues with an AI would not make me feel uncomfortable or insecure.

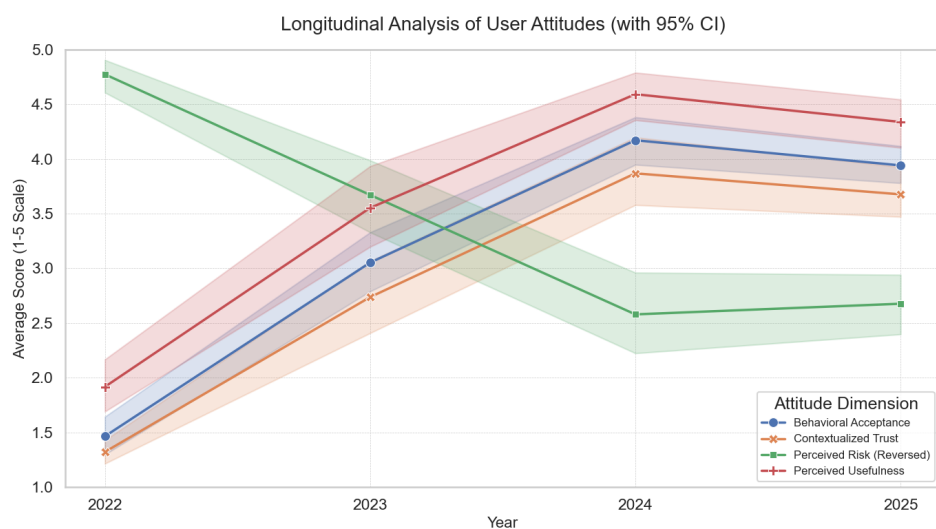


Figure 1: Longitudinal Trajectories of Core Attitude Dimensions with 95% Confidence Intervals (2022–2025).

Table 5

Semi-Structured Interview Protocol for User Attitudes toward Generative AI in Healthcare.

Dimension	No.	Interview Question
Foundational Information	Q1	Gender
	Q2	Age
	Q3	Professional Background
	Q4	Highest Level of Education
	Q5	Please self-rate your familiarity with AI (e.g., ChatGPT, Ernie Bot, etc.).
	Q6	Have you had any experience using AI health assistants or online AI consultation tools?
Introduction & Overall Perception	Q7	When you hear the terms "AI doctor" or "AI health assistant," what is the first image or scenario that comes to mind? Is this feeling generally positive, negative, or complex? In your past experiences seeking health information or medical services, have you ever had a moment where you wished for a "smarter, more convenient tool"?
	Q8	What was that situation like?
	Q9	Compared to traditional methods like searching online, reading health manuals, or consulting a doctor directly, what do you think is the biggest difference or potential advantage of AI in providing healthcare services?
Perceived Usefulness & Application Scenarios	Q10	Imagine the entire process from "feeling slightly unwell" to "visiting a doctor" and then to "recovering at home." In which specific stages could an AI provide help? Please describe what you believe would be the most valuable moment for its intervention. In your opinion, what types of health issues would an AI be particularly good at handling
	Q11	(e.g., common illness consultation, medication reminders, health education)? Conversely, what health issues do you consider to be absolute "red lines" that an AI should not touch?
	Q12	Do you think using an AI would make you feel more in control of your health, or would it make you more anxious due to information overload? Why?
	Q13	Do you expect the health advice provided by an AI to be completely standardized or highly personalized? To what extent would you be willing to share your personal lifestyle, dietary habits, or even genetic information to receive personalized advice?
	Q14	Beyond being free of charge, what other conditions would need to be met (e.g., endorsement by a reputable hospital, a clear user interface, a guarantee of privacy and security) for you to seriously consider trying an AI health assistant?
	Q15	Under what circumstances would you consider consulting an AI to be a "substitute" for a non-essential visit to the doctor? Conversely, under what circumstances would consulting an AI make you feel a more urgent need to see a human doctor?
Personal Use Intention & Acceptance Boundaries	Q16	On one hand, an AI could offer 24/7 emotional support; on the other hand, confiding deep emotions in a machine can be unsettling. How do you personally weigh this "convenience" against the "discomfort"?
	Q17	Do you prefer to view the AI as a one-time information query tool, or as a long-term "digital companion" that can track your health trajectory?
		Why?

Table 6

Semi-Structured Interview Protocol for User Attitudes toward Generative AI in Healthcare (continued).

Dimension	No.	Interview Question
Trust Construction & Collapse	Q18	What does a medical AI need to demonstrate to you to win your initial trust? (e.g., transparent information sources, expert certifications, positive user reviews, or its clear logical reasoning process?)
	Q19	If an AI cannot explain the detailed reasoning behind its judgment like a human doctor can (the "black box" problem), to what extent would this affect your trust in it?
		If you discovered that an AI provided inaccurate information, how would you react?
	Q20	Would this single experience permanently destroy your trust, or would you be willing to try again if it had a good error-correction and apology mechanism?
	Q21	When an AI's advice contradicts information you found online, who are you more inclined to believe? When an AI's advice contradicts that of a community doctor, how would you decide?
Risk Perception & Ethical Considerations	Q22	Of the potential risks associated with AI in healthcare (e.g., misinformation, privacy breaches, algorithmic bias, depersonalization of the doctor-patient relationship, skill degradation in human doctors), which one makes you the most uneasy? Why?
		Regarding data privacy, what are your most specific concerns?
	Q23	Are you afraid of receiving targeted pharmaceutical ads, your data being sold to insurance companies and affecting your premiums, or your personal information being stolen by hackers for fraud?
	Q24	Imagine a scenario where a patient follows an AI's diagnosis and treatment advice perfectly, but their condition worsens. In your opinion, who should be held accountable: the AI developer, the hospital using it, the supervising doctor, or the patient themselves?
	Q25	Are you concerned that advanced AI medical tools like these might exacerbate health inequalities between different population groups (e.g., young vs. old, wealthy vs. poor)?
Role Outlook & Emotional Summary	Q26	If you were to use a metaphor to describe the ideal role of AI in future healthcare, what would it be (e.g., a tool like a "stethoscope," an assistant like a "nurse," a knowledge base like an "encyclopedia," or a collaborator like a "second opinion doctor")? Why that metaphor?
	Q27	If our conversation today had a central theme, would you describe it as "cautious optimism," "deep concern," "practical welcome," or something else?
		Please summarize your core feelings about the integration of AI into healthcare in a few sentences.
	Q28	Beyond everything we have discussed, are there any other thoughts on this topic that you feel are very important, but I have not asked about?

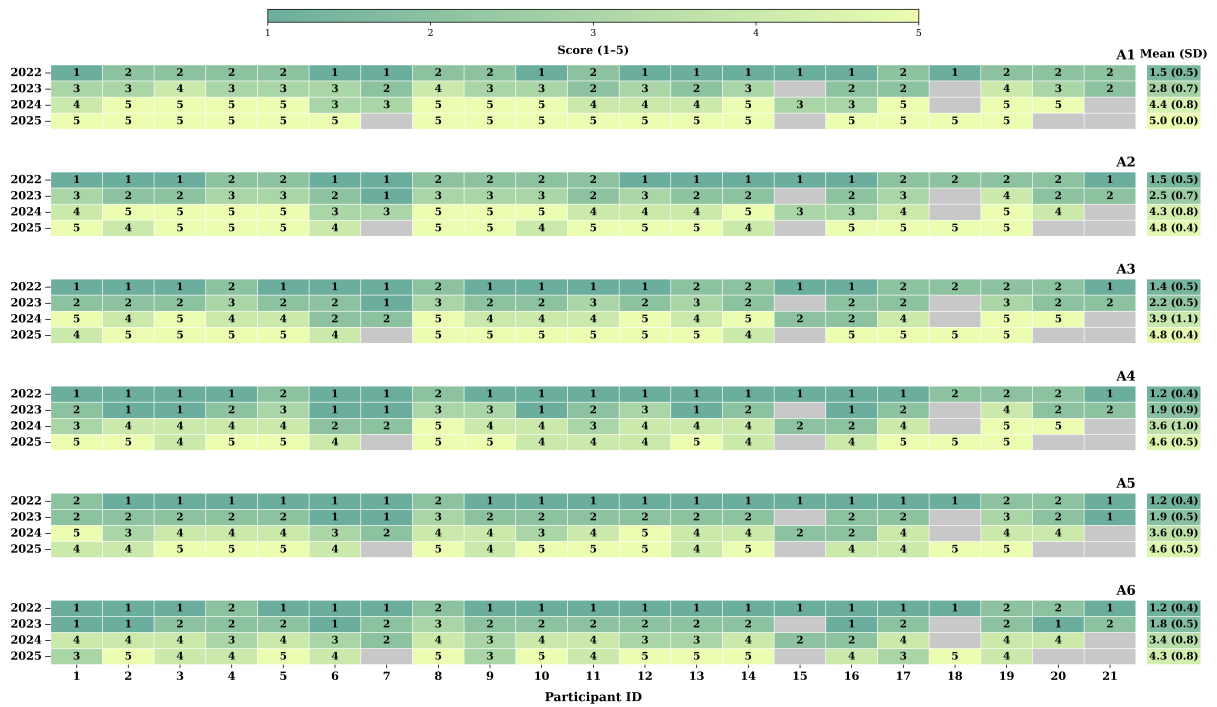


Figure 2: Heatmap of Individual Participant AI Literacy Scores Across Four Waves (2022-2025)

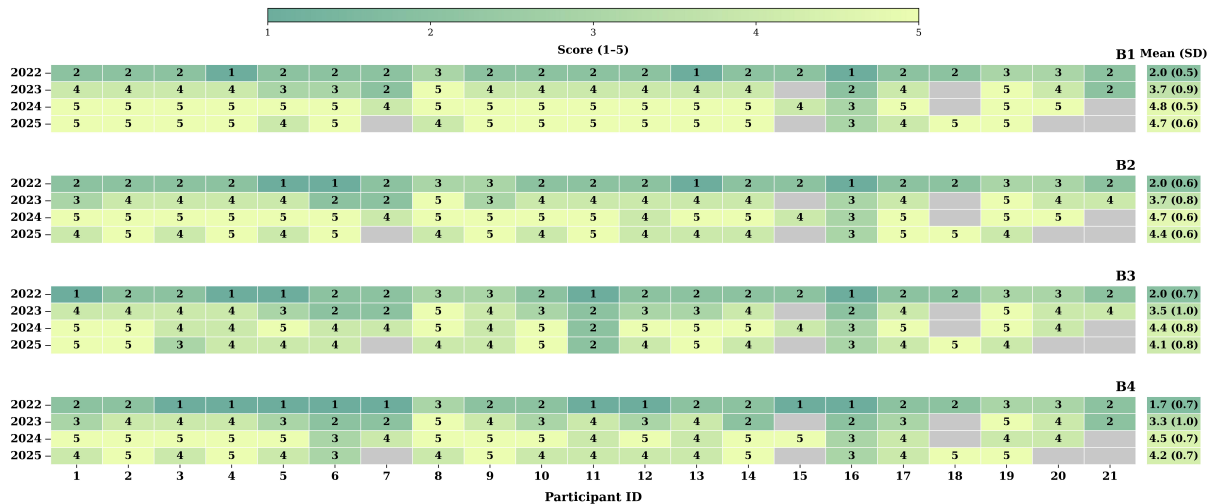


Figure 3: Heatmap of Individual Participant Perceived Usefulness Scores Across Four Waves (2022-2025)

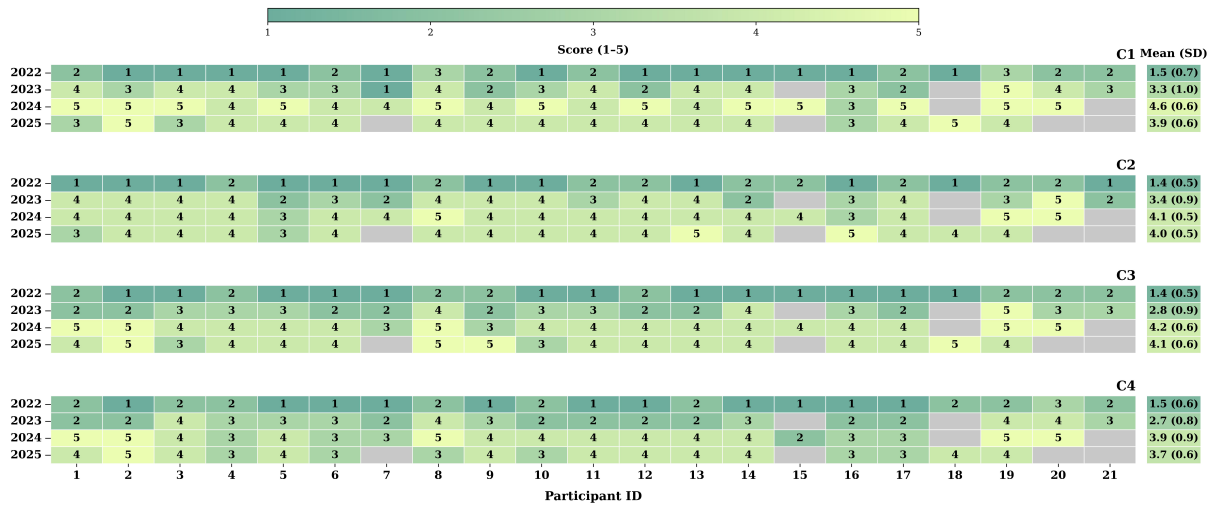


Figure 4: Heatmap of Individual Participant Behavioral Acceptance Scores Across Four Waves (2022-2025)

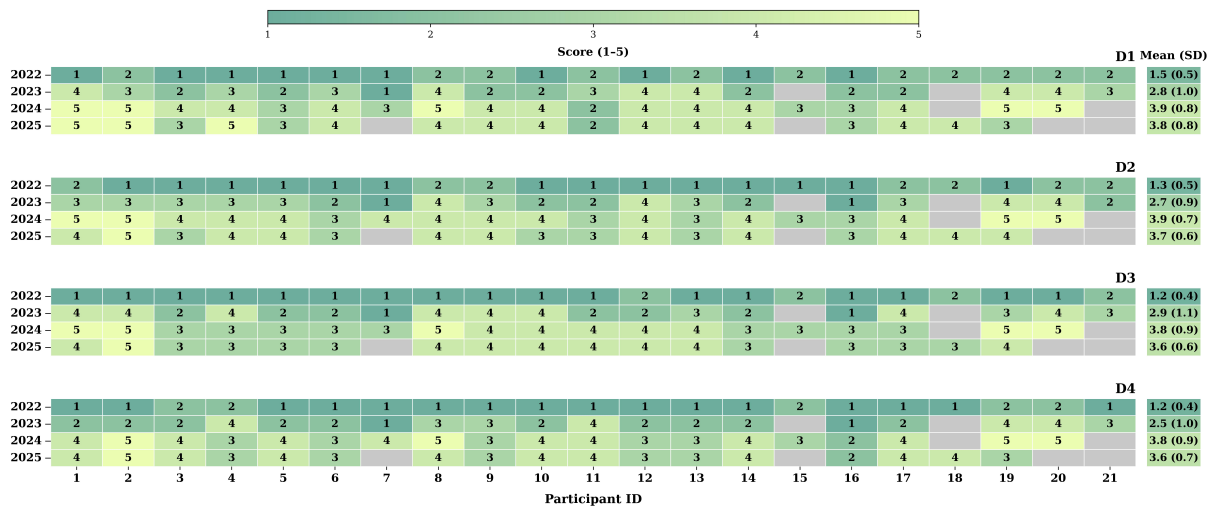


Figure 5: Heatmap of Individual Participant Contextualized Trust Scores Across Four Waves (2022-2025)

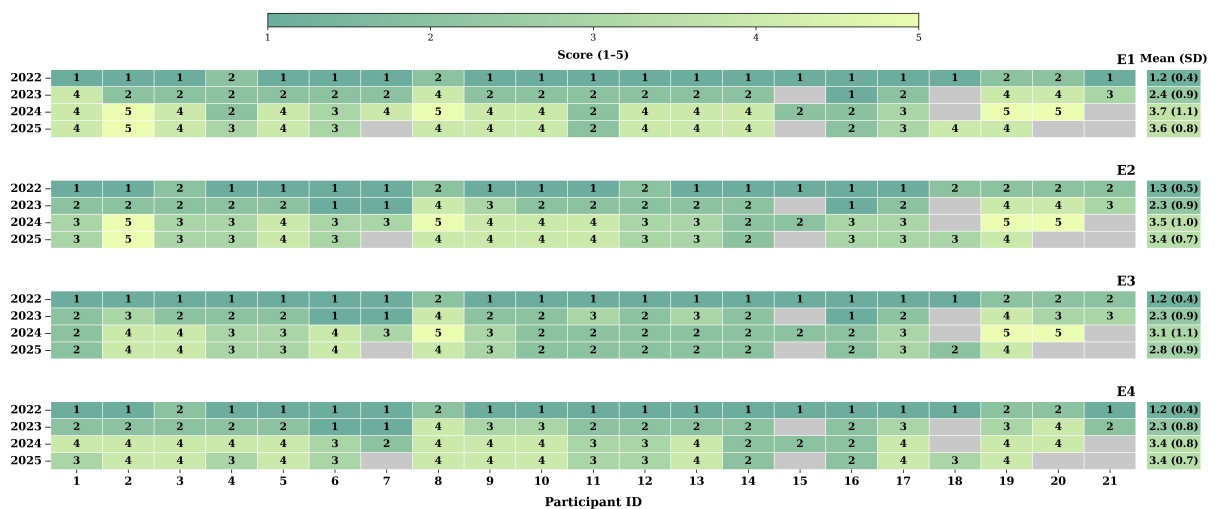


Figure 6: Heatmap of Individual Participant Perceived Risk Scores Across Four Waves (2022-2025)