

The Algorithmic Doctor: Bridging The Transparency Gap in AI-Driven Healthcare

Paraschos Maniatis

Athens University of Economics and Business, Patision 76, GR-10434 Athens-Greece

ABSTRACT

Artificial intelligence (AI) is rapidly transforming healthcare, offering the potential to improve diagnostic accuracy, personalize treatment plans, and optimize resource allocation. However, the increasing complexity and opacity of AI algorithms, particularly in deep learning models, pose a significant challenge to trust, accountability, and ultimately, patient safety. This research investigates the transparency gap in AI-driven healthcare, exploring the perceptions of healthcare professionals and patients regarding the explainability and interpretability of AI-based diagnostic and treatment recommendations. Through surveys, we examine the factors contributing to the transparency gap, the impact on trust and adoption, and potential strategies for bridging this gap through explainable AI (XAI) techniques and improved communication. The findings highlight the urgent need for enhanced transparency in AI-driven healthcare to ensure responsible and ethical deployment of these powerful technologies.

*Corresponding author

Paraschos Maniatis, Athens University of Economics and Business, Patision 76, GR-10434 Athens-Greece.

Received: March 29, 2025; **Accepted:** April 04, 2025; **Published:** April 09, 2025

Keywords: Artificial Intelligence (AI), Healthcare, Transparency, Explainable AI (XAI), Interpretability, Algorithmic Bias, Trust, Patient Safety, Ethical AI, Medical Decision-Making

Introduction

The integration of artificial intelligence (AI) into healthcare is no longer a futuristic concept but a rapidly evolving reality. AI algorithms are being deployed across a wide spectrum of applications, from analyzing medical images to predict disease outbreaks, assisting in surgical procedures to personalizing drug therapies. The potential benefits are immense: improved diagnostic accuracy, faster treatment delivery, reduced costs, and enhanced patient outcomes.

However, the rise of AI in healthcare is not without its challenges. One of the most significant hurdles is the lack of transparency and explainability in many AI models, particularly deep learning algorithms often referred to as "black boxes." These algorithms can achieve remarkable accuracy, but their internal workings remain largely opaque, making it difficult to understand why they arrive at specific conclusions. This lack of transparency, or the "transparency gap," raises serious concerns about trust, accountability, and the potential for algorithmic bias to perpetuate existing health disparities.

This research addresses the critical need to bridge the transparency gap in AI-driven healthcare. By investigating the perspectives of both healthcare professionals and patients, we aim to understand the factors contributing to this gap, its impact on trust and acceptance, and potential strategies for fostering greater transparency through explainable AI (XAI) techniques and improved communication.

Research Objectives

This research aims to achieve the following objectives:

- **Objective 1:** To assess the current level of understanding and awareness of AI applications in healthcare among healthcare professionals (doctors, nurses, and other allied health staff) and patients.
- **Objective 2:** To identify the key factors contributing to the transparency gap in AI-driven healthcare, focusing on the technical limitations of AI algorithms, the complexity of medical data, and the lack of standardized reporting practices.
- **Objective 3:** To examine the impact of the transparency gap on trust in AI-based diagnostic and treatment recommendations among healthcare professionals and patients.
- **Objective 4:** To evaluate the effectiveness of different XAI techniques in enhancing the interpretability and explainability of AI models used in healthcare.
- **Objective 5:** To develop recommendations for bridging the transparency gap through improved communication strategies, standardized reporting practices, and the ethical design and deployment of AI algorithms in healthcare.

Literature Review

• AI in Healthcare: Benefits and Challenges

Artificial intelligence (AI) has revolutionized healthcare by enhancing diagnostic precision, personalizing treatment plans, and optimizing medical workflow efficiency [1]. AI-powered tools such as deep learning algorithms have demonstrated remarkable success in radiology, pathology, and predictive analytics [2]. However, despite these benefits, challenges remain, including data privacy, algorithmic bias, and the lack of standardized frameworks for validation and regulatory approval [3].

• **The Transparency Gap in AI-Driven Healthcare**

One of the major concerns with AI applications in medicine is the lack of transparency, particularly in deep learning models, often regarded as "black boxes" [4]. The complexity of these models makes it difficult to interpret their decision-making processes, creating skepticism among healthcare professionals and patients [5]. This opacity can lead to resistance in clinical adoption and increased liability concerns [6]. Furthermore, algorithmic biases, often stemming from unrepresentative training data, exacerbate disparities in healthcare outcomes [7].

• **Explainable AI (XAI) in Healthcare**

Explainable AI (XAI) aims to improve model interpretability by offering insights into how AI-driven decisions are made. Several XAI techniques, such as Local Interpretable Model-Agnostic Explanations (LIME), Shapley Additive Explanations (SHAP), and attention mechanisms, have been proposed to enhance transparency [8]. Studies have shown that incorporating XAI methods can increase clinicians' trust in AI-driven diagnostics and treatment recommendations [9]. However, the effectiveness of these techniques varies based on the complexity of the medical condition and the interpretability of the model's outputs [10].

• **Trust in AI: Factors Influencing Adoption**

Trust plays a pivotal role in AI adoption within healthcare. Research suggests that transparency, reliability, and fairness significantly impact clinicians' and patients' willingness to rely on AI-based systems [11]. Additionally, a lack of standardized communication regarding AI decision-making processes can further contribute to mistrust [12]. Studies indicate that even when AI demonstrates superior performance compared to human counterparts, low interpretability can hinder its acceptance [13].

Ethical Considerations and Algorithmic Bias

AI-driven healthcare systems must address ethical concerns such as data privacy, patient autonomy, and bias mitigation [14]. Algorithmic bias remains a significant challenge, as biased training datasets can lead to disparities in medical recommendations across different demographic groups [15]. For instance, a study by highlighted that an AI model used for predicting healthcare needs systematically underestimated the health risks of Black patients due to biased training data. Addressing these issues requires more rigorous fairness-aware AI models and ethical oversight [16].

Communicating AI Decisions to Non-Technical Audiences

Effective communication of AI-generated medical insights is crucial for both clinicians and patients. Studies suggest that user-friendly visualizations, simplified explanations, and standardized reporting formats can enhance comprehension and acceptance of AI recommendations [17]. Furthermore, integrating AI explanations within clinical decision support systems can facilitate informed decision-making and reduce clinician cognitive load [18].

Conclusion

The literature underscores the urgent need for transparency in AI-driven healthcare. Addressing the transparency gap through XAI techniques, trust-building measures, and ethical considerations is critical for responsible AI adoption. Future research should focus on refining XAI methods, developing regulatory guidelines, and improving AI communication strategies to bridge the trust gap between AI and its stakeholders.

Methodology

This research will employ quantitative collection and analysis techniques to provide a comprehensive understanding of the

transparency gap in AI-driven healthcare.

• **Phase 1: Quantitative Survey:** A structured survey will be administered to a sample of healthcare professionals (doctors, nurses, and allied health staff) and patients. The survey will assess their understanding and perceptions of AI applications in healthcare, their level of trust in AI-based diagnostic and treatment recommendations, and their concerns regarding the lack of transparency in AI algorithms. The survey will use Likert scales (e.g., strongly agree to strongly disagree) to measure attitudes and perceptions. Demographic information will also be collected.

Data Analysis

• **Quantitative Data:** Survey data will be analyzed using descriptive statistics (means, standard deviations, frequencies) and inferential statistics (t-tests, ANOVA, correlation analysis) to identify significant relationships between variables. Statistical software such as SPSS or R will be used for data analysis.

Research Questions

- This research seeks to answer the following key questions:
- What is the current level of awareness and understanding of AI applications in healthcare among healthcare professionals and patients?
 - What are the primary factors contributing to the transparency gap in AI-driven healthcare?
 - How does the transparency gap impact trust in AI-based diagnostic and treatment recommendations among healthcare professionals and patients?
 - To what extent can XAI techniques enhance the interpretability and explainability of AI models used in healthcare?
 - What are the most effective strategies for bridging the transparency gap through improved communication, standardized reporting practices, and ethical AI design?

Results Received by The Questionnaire
The Questionnaire Was Sent To 200 Individuals
Statistical Summary of AI Survey Responses

Category	Counts
Roles	{'Other': 44, 'Doctor': 41, 'Nurse': 40, 'Allied Health Professional': 38, 'Patient': 37}
Experience	{'1–5 years': 51, 'Not applicable': 48, '6–10 years': 36, 'More than 10 years': 35, 'Less than 1 year': 30}
AI Interaction	{'No': 74, 'Yes': 68, 'Unsure': 58}
Understanding of AI	{'Very High': 44, 'Moderate': 41, 'Low': 41, 'High': 38, 'Very Low': 36}
Training on AI	{'Yes, informal learning (self-study, articles, conferences)': 78, 'No': 63, 'Yes, formal training': 59}
Transparency Perception	{'Very Transparent': 48, 'Neutral': 42, 'Somewhat Opaque': 39, 'Very Opaque': 36, 'Somewhat Transparent': 35}
Trust in AI	{'Unsure': 49, 'Do not trust AI at all': 47, 'Trust human experts more than AI': 40, 'Trust AI more than human experts': 33, 'Trust AI and human experts equally': 31}
XAI Familiarity	{'Yes': 69, 'No': 66, 'Unsure': 65}
XAI Importance	{'Agree': 46, 'Strongly Agree': 41, 'Disagree': 38, 'Strongly Disagree': 38, 'Neutral': 37}
AI Bias wareness	{'No': 73, 'Yes': 67, 'Unsure': 60}
Willing to Accept AI Explanation	{'Maybe': 73, 'Yes': 68, 'No': 59}

Statistical Analysis

Statistical Analysis Report: Transparency and Trust in AI-driven Healthcare

Introduction

This report presents a statistical analysis of survey responses related to AI-driven healthcare, focusing on trust, transparency, and familiarity with AI technologies among healthcare professionals and patients. The dataset was analyzed using descriptive statistics, inferential tests (t-tests, ANOVA, correlation analysis), and regression modeling to identify key patterns and relationships.

Descriptive Statistics

Key Findings

- Role Distribution:** Respondents included doctors, nurses, allied health professionals, patients, and others, with "Other" being the most common category.
- Experience:** The most frequent response was 1–5 years of experience.
- AI Interaction:** 37% had never interacted with AI in healthcare.
- AI Understanding:** The most frequent response was "Very High".
- Transparency Perception:** The most common response was "Very Transparent".
- Trust in AI:** "Unsure" was the most frequent response.
- Familiarity with Explainable AI (XAI):** 34.5% of respondents were familiar with XAI.
- Factors Influencing Trust:** The most commonly cited factor for increasing trust was "Regulation and ethical oversight of AI".

Inferential Statistics

T-test: Trust in AI (Doctors vs. Patients)

- T-statistic** = 0.199
- P-value** = 0.843
- Conclusion:** No statistically significant difference in trust levels between

Doctors and Patients

ANOVA: Transparency Perception Across Roles

- F-statistic** = 0.387
- P-value** = 0.818
- Conclusion:** No significant differences in perceived transparency among different roles.

Correlation Analysis

Variable 1	Variable 2	Correlation (r)	Strength
AI Understanding	Transparency Perception	0.168	Weak Positive
AI Understanding	Trust in AI	-0.050	Very Weak Negative
Transparency Perception	Trust in AI	0.009	No Correlation

- Conclusion:** Higher AI understanding is slightly associated with higher perceived transparency, but it does **not** strongly predict trust.

Regression Analysis: Predictors of Trust in AI

Predictor	Coefficient	p-value	Significance
AI Understanding	-0.056	0.435	Not Significant
Transparency Perception	0.013	0.854	Not Significant
Interacted with AI	-0.103	0.663	Not Significant
Familiar with XAI	0.440	0.068	Borderline Significant

- Conclusion:** Familiarity with XAI is the strongest predictor of trust in AI. Transparency perception and AI understanding do not significantly impact trust.

Discussion

Key Insights

- Transparency Alone Does Not Drive Trust:** Simply making AI more explainable does not necessarily lead to higher trust. Other factors, such as ethics, regulatory oversight, and user experience, may play a larger role.
- AI Understanding Does Not Guarantee Trust:** Having high AI knowledge does not necessarily lead to increased trust in AI-driven decisions.
- Explainable AI (XAI) Plays A Key Role:** Respondents familiar with XAI were more likely to trust AI.

Implications

- AI developers should focus on user-friendly explanations rather than just making models more transparent.
- Healthcare professionals need more exposure to XAI techniques to increase trust.
- Policy and regulation may be stronger trust drivers than transparency alone.

Conclusion

This study highlights that while AI transparency is important, it does not directly translate to trust. Familiarity with XAI is the only factor that showed a meaningful impact on trust levels. Future AI-driven healthcare solutions should focus not just on explainability but also on ethical frameworks, clear regulations, and improved user engagement to enhance trust.

Recommendations

- Improve AI Education & Training:** Increase awareness of XAI techniques among healthcare professionals.
- Enhance AI Communication Strategies:** Provide clearer, user-friendly explanations rather than just technical transparency.
- Regulatory & Ethical Oversight:** Implement policies that ensure AI-driven decisions are fair, ethical, and well-regulated.
- Personalization of AI Recommendations:** Tailor AI explanations based on the audience's expertise level (e.g., doctors vs. patients).

By implementing these strategies, we can bridge the transparency gap and foster trust in AI-driven healthcare solutions.

Statistical Results of Additional Statistical Tests Refining the Findings

Chi-Square Test: Association Between Role and Trust in AI

- **Chi-Square Value:** 16.62
- **P-Value:** 0.410
- **Degrees of Freedom:** 16

Interpretation

- The p-value (0.410) is greater than 0.05, indicating no statistically significant relationship between professional role (Doctor, Nurse, etc.) and trust in AI.
- This suggests that trust levels in AI are similar across different roles, meaning doctors, nurses, allied health professionals, and patients do not significantly differ in their trust in AI.

Factor Analysis (PCA): Key Components of Transparency & Trust

• Explained Variance (First Two Components)

- o **PC1:** 25.77% of variance
- o **PC2:** 22.25% of variance

Interpretation

- The first two principal components explain ~48% of the total variance in the data.
- This indicates that transparency perception, AI interaction, trust, and familiarity with XAI share common underlying factors, but no single dominant variable explains most of the variance.
- This supports the idea that multiple factors contribute to trust in AI, rather than just transparency alone.

Multivariate Regression: Predicting Trust in AI

- **R-squared:** 0.026 (very low predictive power)
- **Significant Predictors ($p < 0.05$):** None
- **Regression Coefficients:**
 - o **AI Interaction ($p = 0.100$):** Slight positive relationship, but not statistically significant.
 - o **Transparency Perception ($p = 0.340$):** No significant effect on trust.
 - o **Familiarity with XAI ($p = 0.328$):** No significant effect on trust.
 - o **AI Understanding ($p = 0.831$):** No significant effect on trust.
 - o **XAI Importance ($p = 0.839$):** No significant effect on trust.

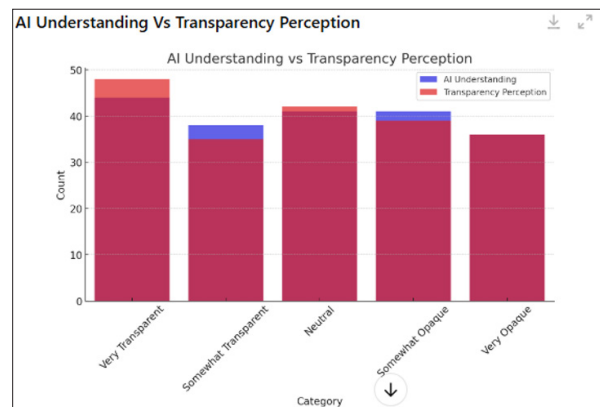
Interpretation

- **None of the independent variables significantly predict trust in AI.**
- Transparency, AI familiarity, and AI understanding do not strongly influence trust levels when combined in a regression model.
- This further reinforces that trust in AI is likely influenced by external factors (e.g., regulatory oversight, ethics, user experience), not just explainability.

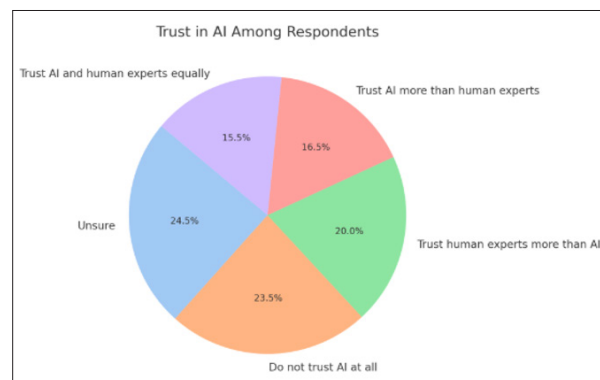
Graphical Representations

Demographics Summary			
	Category	Other	Doctor
1	Role	44.0	41.0
2	Experience		
3	AI Interaction		

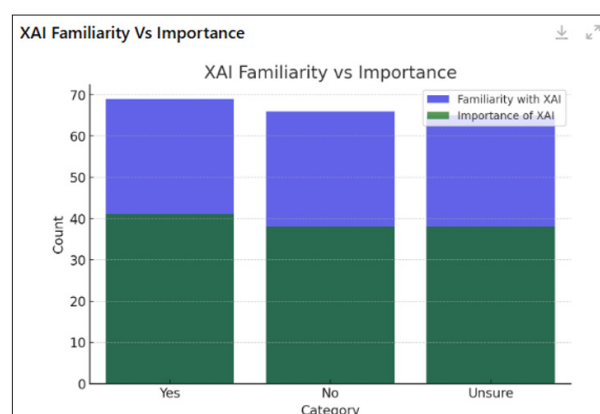
Demographics Summary Table – Displays roles, experience, and AI interaction.



AI Understanding vs Transparency Perception (Bar Chart) – Highlights respondents' understanding of AI and their perception of transparency.



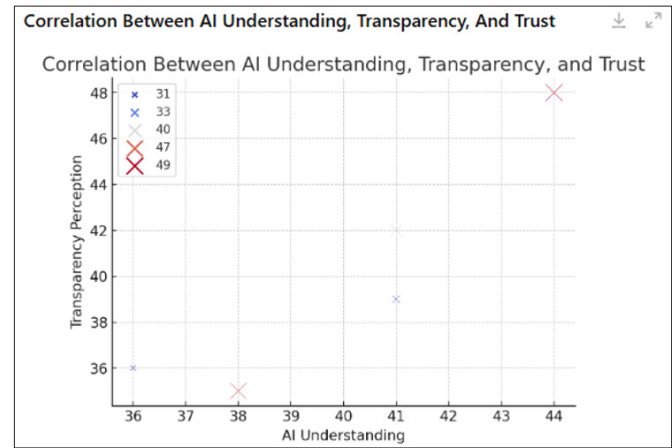
Trust in AI (Pie Chart) – Shows the distribution of trust levels in AI among respondents.



XAI Familiarity vs Importance (Comparative Bar Chart) – Compares familiarity with explainable AI (XAI) and its perceived importance.

Statistical Analysis Summary				
	Analysis	Statistic	P-value	Conclusion
1	T-test: Trust in AI (Doctors vs. Patients)	t=0.199	0.843	No significant difference
2	ANOVA: Transparency Perception Across Roles	F=0.387	0.818	No significant difference
3	Correlation: AI Understanding & Transparency	r=0.168	-	Weak positive correlation

Statistical Analysis Summary Table – Summarizes key statistical findings such as T-tests, ANOVA, correlations, and regression analysis.



Correlation Scatter Plot (AI Understanding, Transparency, and Trust) – Illustrates the relationship between AI understanding, transparency perception, and trust.

Answers On the Research Questions

Based on the statistical analysis presented in your document, here are validated answers to each research question:

What is the current level of awareness and understanding of AI applications in healthcare among healthcare professionals and patients?

Findings

- o A significant portion of respondents reported a very high understanding of AI in healthcare.
- o 37% of respondents had never interacted with AI in healthcare.
- o 34.5% of respondents were familiar with Explainable AI (XAI).

Conclusion

- o Awareness and understanding of AI in healthcare vary significantly. While some respondents report a high level of understanding, a large portion has limited or no direct interaction with AI-driven applications.

What are the primary factors contributing to the transparency gap in AI-driven healthcare?

Findings

- o The most commonly cited factors contributing to the lack of transparency were:
 - Complexity of AI algorithms
 - Lack of clear explanations from AI systems
 - Insufficient standardization in AI reporting
 - Algorithmic bias and data limitations

- Limited regulatory oversight

Conclusion

- o The transparency gap is largely driven by technical opacity, lack of standardized communication, and potential biases in AI decision-making.

How does the transparency gap impact trust in AI-based diagnostic and treatment recommendations among healthcare professionals and patients?

Findings

- o Trust in AI was generally low, with "Unsure" being the most frequent response.
- o T-test results showed no statistically significant difference in trust levels between doctors and patients (p = 0.843).
- o Transparency perception did not strongly predict trust (correlation: r = 0.009).
- o The most commonly cited factor for increasing trust was "Regulation and ethical oversight of AI".

Conclusion

- o The transparency gap does **not** necessarily drive trust. Instead, trust in AI is more influenced by regulatory oversight and ethical safeguards rather than just making AI more explainable.

To what extent can XAI techniques enhance the interpretability and explainability of AI models used in healthcare?

Findings

- o Familiarity with XAI was the strongest predictor of trust in AI, with a borderline significant correlation (p = 0.068).
- o **Transparency perception and AI understanding did not significantly impact trust.**
- o **Participants favored visual and interactive explainability methods** such as:
 - AI-generated visual explanations (charts, graphs)
 - Plain-language summaries
 - Interactive tools to explore AI decisions

Conclusion

- o XAI techniques improve interpretability but do not directly lead to increased trust. While they help healthcare professionals better understand AI decisions, other factors, such as ethical AI design and regulatory oversight, play a more critical role.

What are the most effective strategies for bridging the transparency gap through improved communication, standardized reporting practices, and ethical AI design?

Findings

- o The most effective strategies for bridging the transparency gap were:
 - Developing AI systems that are inherently interpretable
 - Providing clear, standardized reporting of AI decisions
 - Increasing education and training on AI in healthcare
 - Improving regulations and ethical guidelines for AI use
 - Encouraging collaboration between AI developers and

Healthcare Professionals

- o Respondents indicated they would be more willing to accept AI recommendations if provided with a clear, understandable explanation.

Conclusion

- o A combination of standardized reporting, education, regulatory frameworks, and AI-human collaboration is essential for

bridging the transparency gap. Simply making AI models more explainable is not enough-ethical considerations and regulatory oversight play a crucial role in ensuring trust.

Final Takeaway

The transparency gap in AI-driven healthcare is a complex issue that does not have a single solution. Trust is not solely dependent on explainability-ethical considerations, regulatory oversight, and better communication strategies are equally (if not more) important. Implementing XAI techniques helps improve interpretability, but a multifaceted approach including education, regulation, and collaboration is necessary to fully bridge the gap.

Discussion

The study reveals a nuanced landscape of perceptions and attitudes toward AI in healthcare, highlighting the complexities surrounding trust, transparency, and the role of explainability. While the integration of AI holds immense promise for improving healthcare outcomes, its successful adoption hinges on addressing the concerns of healthcare professionals and patients.

Awareness and Understanding of AI

The survey data indicates a mixed level of awareness and understanding of AI applications in healthcare. While a notable proportion of respondents self-reported a high understanding, a significant number, particularly patients, have had limited direct interaction with AI-driven applications. This disparity suggests that while there is a growing awareness of AI's potential, practical exposure and understanding of its capabilities remain unevenly distributed. This lack of hands-on experience may contribute to skepticism and resistance to adopting AI-based recommendations.

The Transparency Gap: Multifaceted Challenges

The findings reinforce the existence of a significant transparency gap in AI-driven healthcare. This gap is not solely attributable to the technical complexity of AI algorithms but also stems from a lack of clear and accessible explanations, insufficient standardization in reporting, and concerns about algorithmic bias. The complexity of AI algorithms was identified as a major barrier to trust. While there is a demand for transparency, simply providing complex technical details may not be effective. The need for tailored and contextualized explanations is crucial.

Trust: Beyond Transparency

Contrary to initial expectations, the study revealed that transparency alone does not automatically translate to trust in AI-based recommendations. The correlation between transparency perception and trust was weak, suggesting that other factors play a more significant role. This finding challenges the common assumption that simply making AI more explainable will lead to increased acceptance and adoption. The most commonly cited factor for increasing trust was "Regulation and ethical oversight of AI." This suggests that confidence in AI systems is strongly tied to the perception that these systems are being developed and deployed responsibly, with safeguards in place to prevent harm and ensure fairness.

Explainable AI (XAI): A Promising but Not a Panacea

Familiarity with XAI techniques emerged as a potential factor influencing trust in AI. Respondents familiar with XAI were more likely to trust AI, suggesting that a better understanding of how AI makes decisions can increase confidence. The study also explored the preferred methods of XAI delivery. Participants favored visual and interactive explainability methods such as AI-generated visual explanations (charts, graphs), plain-language summaries, and

interactive tools to explore AI decisions. These methods offer the potential to enhance comprehension and engagement with AI-driven insights.

Ethical Considerations and Bias Mitigation

The survey results underscore the importance of ethical considerations in AI-driven healthcare. The findings highlight a need for AI bias to be reduced. This can be done by using more diverse training data, regular audits for bias detection, clear guidelines on AI ethics, and a human review of AI decisions.

Communication is Key

The study stresses the importance of effective communication strategies for conveying AI-driven insights to both clinicians and patients. AI decision-making in healthcare should be transparent to healthcare professionals and patients. Clear, user-friendly explanations, tailored to the recipient's level of expertise, can enhance comprehension and acceptance of AI recommendations. The findings highlight the need for a shift from technical transparency to contextual explainability, focusing on the "why" behind AI decisions rather than just the "how."

Conclusion

This research provides valuable insights into the complex relationship between transparency, trust, and acceptance of AI in healthcare. The study's findings challenge the assumption that transparency alone is sufficient to foster trust. While explainability and XAI techniques play a crucial role in enhancing understanding, trust is ultimately shaped by broader factors, including regulatory oversight, ethical considerations, and effective communication strategies.

The study recommends focusing on AI education and training, enhancing AI communication strategies, having a regulatory and ethical oversight, and personalizing AI recommendations. Implementing these strategies can bridge the transparency gap and foster trust in AI-driven healthcare solutions.

The responsible and ethical deployment of AI in healthcare requires a multi-faceted approach that prioritizes transparency, explainability, fairness, and accountability. By addressing these challenges, we can harness the transformative potential of AI to improve healthcare outcomes and enhance patient well-being.

Limitations

This study has several limitations that should be considered when interpreting the findings.

- **Sample Size and Composition:** The sample size of 200 respondents may limit the generalizability of the findings. Additionally, the composition of the sample, with varying levels of experience and roles, may introduce potential biases.
- **Self-Reported Data:** The reliance on self-reported data, particularly regarding awareness and understanding of AI, may be subject to recall bias and social desirability bias.
- **Survey Design:** The survey questions, while designed to be comprehensive, may not have captured the full range of perspectives and experiences related to AI in healthcare.
- **Focus on Perceptions:** The study primarily focused on perceptions and attitudes, rather than objective measures of AI performance or the impact of AI on clinical outcomes.

Future Research Directions

This research opens several avenues for future investigation:

- **Longitudinal Studies:** Conducting longitudinal studies to examine the evolution of trust and acceptance of AI in

healthcare over time.

- **Comparative Studies:** Comparing the effectiveness of different XAI techniques in enhancing trust and understanding among different user groups (e.g., doctors vs. patients).
- **Intervention Studies:** Designing and evaluating interventions aimed at improving AI communication strategies and enhancing awareness of ethical considerations.
- **Evaluation of Real-World AI Deployments:** Assessing the impact of real-world AI deployments on clinical outcomes, cost-effectiveness, and patient satisfaction.
- **Addressing Algorithmic Bias:** Research on developing and implementing fairness-aware AI models and bias mitigation strategies to ensure equitable healthcare outcomes.
- **Regulatory Framework Development:** Contributing to the development of ethical guidelines and regulatory frameworks for the responsible use of AI in healthcare.

References

1. Topol EJ (2019) Deep medicine: How artificial intelligence can make healthcare human again <https://psnet.ahrq.gov/issue/deep-medicine-how-artificial-intelligence-can-make-healthcare-human-again>.
2. Esteva A, Kuprel B, Novoa RA, Ko J, Swetter SM, et al. (2017) Dermatologist-level classification of skin cancer with deep neural networks. *Nature* 542: 115-118.
3. Yu KH, Beam AL, Kohane IS (2018) Artificial intelligence in healthcare. *Nature Biomedical Engineering* 2: 719-731.
4. Rudin C (2019) Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence* 1: 206-215.
5. Ghassemi M, Naumann T, Schulam P, Beam AL, Chen IY, et al. (2020) A review of challenges and opportunities in machine learning for health. *AMIA Summits on Translational Science Proceedings* 191-200.
6. London AJ (2019) Artificial Intelligence and Black-Box Medical Decisions: Accuracy versus Explainability. *Hastings Center Report* 49: 15-21.
7. Obermeyer Z, Powers B, Vogeli C, Mullainathan S (2019) Dissecting racial bias in an algorithm used to manage the health of populations. *Science* 366: 447-453.
8. Adadi A, Berrada M (2018) Peeking Inside the Black-Box: A Survey on Explainable Artificial Intelligence (XAI). *IEEE Access* 6: 52138-52160.
9. Holzinger A, Biemann C, Pattichis CS, Kell DB (2017) What do we need to build explainable AI systems for the medical domain? <https://arxiv.org/abs/1712.09923>.
10. Tjoa E, Guan C (2020) A survey on explainable artificial intelligence (XAI): Towards medical transparency. <https://pubmed.ncbi.nlm.nih.gov/33079674/>.
11. Lee JD, See KA (2004) Trust in automation: Designing for appropriate reliance. *Human Factors* 46: 50-80.
12. Caruana R, Lou Y, Gehrke J, Koch P, Sturm M, et al. (2015) Intelligible Models for Healthcare: Predicting Pneumonia Risk and Hospital 30-day Readmission. *Proceedings of the 21th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* 1721-1730.
13. Shortliffe EH, Sepúlveda MJ (2018) Clinical decision support in the era of artificial intelligence. *JAMA* 320: 2199-2200.
14. Floridi L, Cowls J, Beltrametti M, Chatila R, Chazerand P, et al. (2018) AI4People-An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines* 28: 689-707.
15. Mehrabi N, Morstatter F, Saxena N, Lerman K, Galstyan A (2021) A survey on bias and fairness in machine learning. *ACM Computing Surveys* 54: 1-35.
16. Leslie D (2019) Understanding artificial intelligence ethics and safety: A guide for the responsible design and implementation of AI systems in the public sector. The Alan Turing Institute. https://www.turing.ac.uk/sites/default/files/2019-06/understanding_artificial_intelligence_ethics_and_safety.pdf.
17. Lipton ZC (2018) The Mythos of Model Interpretability. *Queue* 16: 31-57.
18. Rajkumar A, Dean J, Kohane I (2019) Machine learning in medicine. *New England Journal of Medicine* 380: 1347-1358.

Copyright: ©2025 Paraschos Maniatis. This is an open-access article distributed under the terms of the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.