

Enhancing Human Trust Through Explainable AI Communication for Nontechnical Stakeholders

Dwi Apriliasari^{1*} , Bintang Nandana Henry² , Alexander Williams³

¹Asosiasi Dosen Indonesia, Indonesia

²Kreatif Desain, Indonesia

³Queensland University, Australia

¹dwi.apriliasari@raharja.info, ²bintang.nandana@raharja.info, ³a87.williams@connect.qut.edu.au

*Corresponding Author

Article Info

Article history:

Submission May 24, 2026

Revised July 7, 2026

Accepted July 23, 2026

Published August 3, 2026

Keywords:

Explainable AI

Human Trust

AI Communication

Human Centered AI

Nontechnical



ABSTRACT

The rapid adoption of Artificial Intelligence (AI) in data driven decision making has increased the complexity of analytical models, creating significant communication barriers between technical experts and nontechnical stakeholders. Limited understanding of AI generated insights often reduces trust, delays decision making, and restricts the effective adoption of AI supported recommendations. As organizations increasingly rely on explainable and human centered AI systems, effective communication has become essential to ensuring that AI outputs are accessible, transparent, and meaningful for diverse stakeholder groups. **This study aims** to identify the key challenges in communicating complex AI model results and to develop practical communication strategies that enhance human trust among nontechnical stakeholders. **A qualitative** research design was employed using case studies, open ended surveys, expert interviews, and focus group discussions involving data scientists and nontechnical decision makers from business organizations. The collected data were analyzed through thematic analysis to identify recurring communication barriers and effective explanatory practices. **The findings** reveal that technical jargon, cognitive overload, and limited contextual explanations are the primary factors reducing stakeholder trust in AI generated insights. Conversely, explainable AI communication supported by intuitive data visualization, contextual storytelling, simplified summaries, and audience centered messaging significantly improves understanding, transparency, and stakeholder confidence across organizational contexts. **The study proposes** a human centered communication framework that strengthens trust in AI assisted decision making while promoting more inclusive and responsible technology adoption. These findings contribute to explainable AI research by demonstrating how effective communication can bridge the gap between technical complexity and human understanding, thereby generating meaningful humanistic impacts in organizational decision making and sustainable organizational innovation.

This is an open access article under the [CC BY 4.0](https://creativecommons.org/licenses/by/4.0/) license.



This is an open-access article under the CC-BY license (<https://creativecommons.org/licenses/by/4.0/>)

©Authors retain all copyrights

1. INTRODUCTION

Artificial Intelligence (AI) has become a transformative technology that reshapes decision making across healthcare, finance, education, manufacturing, and public administration. Advances in machine learning and predictive analytics enable organizations to process vast amounts of data and generate highly accu-

rate recommendations, creating significant opportunities to improve organizational performance and public services. However, the increasing complexity of AI models has intensified concerns regarding transparency, interpretability, and user acceptance, particularly among nontechnical stakeholders. Although sophisticated AI models often achieve superior predictive performance, their decision-making processes remain difficult to understand, creating communication barriers that reduce confidence in AI-assisted decisions [1, 2]. Consequently, Explainable Artificial Intelligence (XAI) and Human-Centered AI have emerged to emphasize not only computational accuracy but also human understanding, trust, and collaboration. This perspective aligns with the United Nations Sustainable Development Goals (SDGs), particularly SDGs 9, which promotes sustainable innovation and resilient infrastructure, and SDGs 16, which advocates transparent, accountable, and trustworthy institutions. Recent studies also demonstrate that transparent AI explanations improve user confidence, reduce uncertainty, and facilitate more informed decision making, highlighting communication as an essential component of responsible AI implementation [3].

Despite substantial progress in explainable AI research, existing studies predominantly focus on algorithmic transparency, visualization techniques, and technical explanation methods, while relatively limited attention has been given to how AI-generated insights are communicated to nontechnical stakeholders [4]. In practice, data scientists and AI practitioners often face challenges in translating sophisticated analytical outputs into language that business managers, policymakers, healthcare professionals, and other decision makers can readily understand and confidently apply [5]. Technical terminology, statistical complexity, and abstract performance metrics frequently create cognitive overload, reducing stakeholders' willingness to adopt AI-supported recommendations even when predictive accuracy is high. Contemporary research increasingly argues that explainability should extend beyond technical model interpretation toward communication strategies that support collaboration between humans and intelligent systems. Moreover, effective communication contributes to inclusive innovation by ensuring that AI-generated knowledge is accessible to individuals with diverse educational, professional, and digital literacy backgrounds, supporting the objectives of SDGs 17 through knowledge sharing and cross-sector collaboration [6, 7].

Although previous studies recognize the importance of explainability and transparency, several important research gaps remain. Most explainable AI research concentrates on computational interpretability methods while providing limited guidance on how AI explanations should be communicated effectively to diverse stakeholder groups. Existing communication-oriented studies often investigate visualization, storytelling, or dashboard usability independently rather than integrating these approaches into a comprehensive human-centered communication framework that simultaneously promotes understanding, trust, and informed decision making [8]. Furthermore, nontechnical stakeholders are frequently treated as passive recipients instead of active participants whose perceptions, digital literacy, and contextual needs influence successful AI adoption. Relatively few studies also connect AI communication strategies with broader humanistic outcomes, including stakeholder empowerment, organizational trust, and inclusive participation. Therefore, a more holistic perspective is required to ensure that AI-generated knowledge remains understandable, meaningful, and actionable for diverse users [9, 10].

To address these gaps, this study investigates how explainable AI communication enhances human trust among nontechnical stakeholders by identifying communication barriers and proposing practical strategies for interpreting complex AI-generated insights. Rather than focusing solely on algorithmic explainability, this research positions communication as the critical bridge between technical intelligence and human understanding [11]. Using qualitative approaches, including case studies, open-ended surveys, expert interviews, and focus group discussions, the study explores how intuitive data visualization, contextual storytelling, simplified summaries, and audience-centered messaging improve transparency, stakeholder engagement, and confidence in AI-assisted decision making.

The primary contribution of this research is the development of a human-centered communication framework that integrates explainable AI principles with communication theory and trust-building perspectives, thereby extending existing literature beyond purely technical interpretations of explainability. The proposed framework contributes theoretically to interdisciplinary research on explainable AI, human-computer interaction, organizational communication, and responsible AI adoption [12], while providing practical guidance for organizations seeking to improve AI communication practices and foster transparent, inclusive, and trustworthy decision-making environments. By strengthening the relationship between explainable AI and human-centered communication, this study advances a broader understanding of how effective communication can enhance trust, facilitate AI adoption, and support the responsible implementation of intelligent systems

across diverse organizational and societal contexts [13, 14].

2. LITERATURE REVIEW

The increasing adoption of XAI has shifted research beyond algorithmic performance toward understanding how AI explanations influence human perception, trust, and decision making. While conventional XAI primarily focuses on computational interpretability, recent studies emphasize that explainability should also address users' cognitive needs, contextual understanding, and decision-making capabilities. This has led to the emergence of Human-Centered Explainable AI (HCXAI), which integrates artificial intelligence, human-computer interaction, psychology, organizational behavior, and communication science [15]. From this perspective, transparency alone is insufficient unless AI explanations are communicated in ways that enable diverse users to understand, evaluate, and confidently apply AI-generated recommendations. Therefore, effective AI communication has become a key element of responsible AI implementation, particularly for nontechnical stakeholders. Furthermore, understanding recent developments in this field is essential for identifying effective communication approaches that enhance transparency, user trust, and AI adoption across different application domains. Table 1 summarizes representative studies published between 2021 and 2026 on explainable AI, human trust, communication strategies, and human-centered AI [16].

Table 1. Literature Synthesis of Explainable AI Communication and Human Trust (2021–2026)

Authors	Research Focus	Method	Key Findings	Research Gap
[17, 18]	Human-Centered	Explainable AI Conceptual Review	Explainability should prioritize user experience rather than algorithmic transparency alone.	No practical communication framework for organizational stakeholders.
[19]	Human Explanations in AI	Conceptual Analysis	Human explanations are contextual, social, and audience dependent.	Limited application in AI-supported organizational decision making.
[20, 21]	Explainable AI and Trust	Literature Review	Explainability positively influences trust, transparency, and AI acceptance.	Communication mechanisms remain insufficiently explored.
[22]	Human-Centered XAI User Studies	Systematic Review	Most XAI evaluations measure trust and usability but lack interdisciplinary integration.	Few studies investigate communication strategies for nontechnical users.
[23]	Responsible Artificial Intelligence	Review	Responsible AI requires transparency, explainability, ethics, and stakeholder engagement.	Limited discussion regarding communication practices.
[24]	Human-Centered Explainable AI	ARIST Review	Human-centered XAI emphasizes user context, actionability, and trust rather than technical explanations alone.	Practical communication frameworks remain underdeveloped.
[25]	Trustworthy AI	Review	Transparency, uncertainty communication, and trust calibration jointly improve trustworthy AI adoption.	Does not specifically address nontechnical stakeholders.

Table 1 demonstrates a clear evolution of explainable AI research from algorithm-centered approaches toward human-centered perspectives. Earlier studies emphasized improving algorithm interpretability and visualization techniques, whereas more recent investigations increasingly recognize that effective explanations

[26]	Human AI Team- ing	Systematic Re- view	User-centered expla- nations strengthen collaboration and trust between humans and AI systems.	Framework focuses on human-AI teaming rather than organiza- tional communication.
[27]	Explainable AI in Healthcare	Empirical Study	Explainability supports trust, collaboration, and informed decisions beyond transparency alone.	Domain-specific find- ings with limited organizational general- ization.
[28]	Human-Centered XAI	Systematic Re- view	Human-centered design improves explainability, usability, and user accep- tance across AI applica- tions.	Communication strate- gies remain fragmented across disciplines.

should account for users' knowledge, cognitive capacity, contextual needs, and trust formation. Furthermore, contemporary research consistently identifies transparency as a necessary condition for responsible AI; however, transparency alone does not automatically generate trust unless supported by effective communication practices that enable users to understand the rationale, uncertainty, and practical implications of AI-generated outcomes [29]. Another important observation is that although human-centered AI has become a rapidly growing research area, existing studies remain fragmented across artificial intelligence, human-computer interaction, psychology, and information systems, resulting in limited integration between explainability, communication, and stakeholder trust. Most existing frameworks continue to emphasize explanation quality from technical or interface perspectives while paying comparatively little attention to communication strategies specifically designed for nontechnical stakeholders operating in organizational environments [30, 31].

Building upon these observations, the present study extends prior literature by positioning explainable AI communication as a strategic mechanism for enhancing human trust rather than merely improving algorithmic transparency. Unlike previous studies that predominantly evaluate explainability through technical performance or user interface design, this research integrates communication theory, human-centered AI principles, and organizational decision-making perspectives into a unified conceptual framework. Specifically, the study investigates how communication strategies including intuitive visualization, contextual storytelling, simplified summaries, and audience-centered messaging facilitate stakeholder understanding, strengthen trust, and support responsible AI adoption among nontechnical decision makers [32]. Based on the literature synthesis, several key dimensions of Human-Centered Explainable AI communication can be identified. These dimensions integrate technical explainability with communication effectiveness and provide the conceptual basis for the proposed research framework. Table 2 summarizes these dimensions [33].

Table 2. Dimensions of Human-Centered Explainable AI Communication

Dimension	Description	Supporting Literature
Transparency	AI explanations clearly describe how recommendations are generated.	[4, 30, 34]
Interpretability	Users can understand the rationale behind AI outputs.	[33, 35, 36]
Communication Clarity	Technical results are translated into simple and understandable language.	[37]
Contextual Relevance	Explanations are linked to users' decision-making context.	[38]

Table 2 synthesizes the principal dimensions of Human-Centered Explainable AI communication derived from the literature reviewed in Table 1. The synthesis demonstrates that explainability is no longer viewed

Audience Adaptation	Explanations are tailored to different stakeholder backgrounds.	[39]
Human Trust	Effective communication increases confidence in AI-assisted decisions.	[3, 13]

solely as a technical characteristic of AI systems but also as a communication process that enables stakeholders to understand, evaluate, and confidently utilize AI-generated recommendations. The identified dimensions including transparency, interpretability, communication clarity, contextual relevance, audience adaptation, and human trust collectively emphasize the importance of aligning AI explanations with users’ cognitive characteristics and decision-making contexts [40]. These dimensions therefore serve as the conceptual basis for developing the research framework proposed in the next section.

2.1. Proposed Research Framework

The synthesis of previous studies indicates that explainable AI communication extends beyond algorithmic transparency by emphasizing how AI generated information is interpreted, understood, and trusted by human users. Although prior research highlights the importance of explainability, relatively few studies integrate communication strategies, stakeholder understanding, and trust development into a comprehensive human centered framework [41]. Therefore, this study proposes a research framework that conceptualizes explainable AI communication as a multidimensional process through which communication effectiveness strengthens stakeholder understanding, enhances human trust, and supports responsible AI assisted decision making. The framework is derived from the literature synthesis and serves as the conceptual foundation for the research methodology and data analysis [42, 43].



Figure 1. Research Framework

Figure 1 illustrates the proposed research framework developed in this study. The framework positions Explainable AI Communication as the primary construct consisting of four complementary communication strategies, namely intuitive data visualization, contextual storytelling, simplified summaries, and audience centered messaging [44]. These communication practices collectively improve communication effectiveness by reducing cognitive complexity and enabling nontechnical stakeholders to interpret AI generated insights more accurately.

Improved communication effectiveness subsequently facilitates greater stakeholder understanding regarding the rationale, limitations, and practical implications of AI recommendations. As understanding in-

creases, stakeholders become more confident in evaluating AI supported decisions, thereby strengthening human trust toward intelligent systems. This process is consistent with the principles of Human Centered AI, which emphasize that successful AI implementation depends not only on computational performance but also on meaningful human interaction and understandable explanations [45].

Finally, enhanced human trust contributes to more responsible AI assisted decision making and generates broader humanistic impacts, including improved organizational transparency, inclusive technology adoption, stronger stakeholder confidence, and better evidence based decisions. Therefore, the proposed framework demonstrates that communication functions as the critical mechanism connecting technical explainability with human centered outcomes, providing the conceptual basis for the qualitative investigation conducted in this study [46, 47].

3. METODOLOGY

3.1. Research Design

This study employed a qualitative research design to investigate how explainable AI communication enhances human trust among nontechnical stakeholders in organizational decision-making contexts. A qualitative approach was considered appropriate because it enables an in-depth exploration of participants' experiences, perceptions, and communication practices when interpreting AI-generated insights. The study involved AI practitioners, data scientists, managers, business analysts, and other nontechnical decision makers selected through purposive sampling based on their experience in developing, communicating, or utilizing AI-assisted recommendations [48, 49]. Data were collected through semi-structured interviews, focus group discussions, and open-ended questionnaires to capture diverse perspectives regarding communication challenges, trust formation, and effective explainable AI communication practices. To strengthen the credibility of the findings, methodological triangulation was applied by integrating evidence from multiple data collection methods and participant groups, thereby providing a comprehensive understanding of the phenomenon under investigation [50].

3.2. Data Analysis

The collected data were analyzed using thematic analysis following the six-phase framework proposed by Braun and Clarke, including data familiarization, initial coding, theme generation, theme refinement, theme definition, and report development. Interview transcripts, focus group discussion records, and questionnaire responses were systematically coded to identify recurring patterns related to explainable AI communication, stakeholder understanding, communication barriers, and human trust [36]. Similar codes were subsequently grouped into broader themes through iterative comparison across all data sources to ensure conceptual consistency. The resulting themes were then interpreted in relation to the proposed research framework, enabling the identification of communication strategies that contribute to improved stakeholder understanding, enhanced human trust, and responsible AI-assisted decision making [51, 52].

4. RESULT AND DISCUSSION

4.1. Communication Challenges in Explainable AI

The qualitative findings indicate that one of the most significant challenges in implementing XAI is the communication gap between technical experts and nontechnical stakeholders. Across interviews, focus group discussions, and open-ended questionnaires, participants consistently reported difficulties in interpreting AI-generated outputs due to the extensive use of technical terminology, statistical metrics, and algorithm-specific explanations. Although AI systems are capable of producing highly accurate predictions and recommendations, these outputs often remain inaccessible to decision makers without technical expertise. Consequently, stakeholders frequently experience uncertainty regarding the rationale behind AI-generated recommendations, reducing their confidence in utilizing AI-assisted insights during organizational decision-making processes [53].

Another recurring issue identified in this study relates to cognitive overload. Participants explained that AI outputs are commonly presented through complex dashboards, detailed performance indicators, or numerical evaluation metrics that require specialized knowledge to interpret correctly. Rather than facilitating decision making, excessive technical information often complicates stakeholders' understanding and shifts their reliance back to intuition or expert opinion. These findings suggest that explainability should not be viewed

solely as an algorithmic characteristic but also as a communication process that transforms complex analytical outputs into understandable and actionable knowledge.

The findings further reveal that communication quality directly influences stakeholders’ perceptions of AI reliability and transparency. Participants indicated that when explanations were presented using simple language, practical examples, and relevant organizational contexts, they became more willing to consider AI-generated recommendations as credible decision-support tools. Therefore, communication emerges as a critical factor connecting technical explainability with human understanding, highlighting the need for communication strategies specifically designed for nontechnical stakeholders.

4.2. Human Centered Communication Strategies

Following the identification of communication barriers, this study explored the communication practices that participants considered most effective for improving their understanding of AI-generated insights. Despite differences in organizational roles and levels of technical expertise, participants consistently emphasized four communication strategies that enhanced explanation clarity, reduced cognitive complexity, and strengthened confidence in AI-assisted decision making. These recurring themes represent the principal findings of this research and are summarized in Table 3.

Table 3. Human Centered Communication Strategies Identified in the Study

Communication Strategy	Description	Contribution to Human Trust
Intuitive Data Visualization	Presenting AI outputs through simple charts, dashboards, icons, and visual summaries that are easy to interpret	Reduces cognitive complexity and improves comprehension
Contextual Storytelling	Explaining AI-generated insights using practical organizational scenarios and real-world examples	Enhances relevance and stakeholder engagement
Simplified Summary	Translating technical findings into concise, understandable, and actionable messages	Improves accessibility and minimizes misunderstanding
Audience Centered Messaging	Tailoring explanations according to stakeholders’ technical knowledge, responsibilities, and decision contexts	Strengthens confidence and promotes trust in AI-supported recommendations

Table 3 demonstrates that effective explainable AI communication requires more than transparent algorithms or technically accurate explanations. Instead, participants highlighted the importance of presenting AI-generated insights in ways that align with human cognitive processes and decision-making needs. Intuitive data visualization enables stakeholders to quickly identify key information without becoming overwhelmed by excessive analytical detail. Similarly, contextual storytelling bridges the gap between abstract AI outputs and practical organizational situations, allowing users to understand not only what the AI recommends but also why the recommendation is relevant.

Furthermore, simplified summaries help stakeholders focus on the most important findings while reducing the risk of misinterpretation caused by technical jargon. Audience-centered messaging further reinforces communication effectiveness by recognizing that different stakeholder groups possess varying levels of AI literacy, professional expertise, and information requirements. Collectively, these strategies support the principles of Human-Centered AI by demonstrating that communication quality is a fundamental prerequisite for improving stakeholder understanding and fostering trust in AI-assisted decision-making.

4.3. Practical Implications for Responsible AI

The findings of this study indicate that effective explainable AI communication generates practical implications extending beyond technical system implementation. Organizations should recognize communication as a strategic component of AI deployment by ensuring that analytical outputs are presented in ways that are understandable, transparent, and relevant to stakeholders with different levels of technical expertise. Likewise,

AI developers should design explanation mechanisms that prioritize human understanding rather than focusing exclusively on algorithmic accuracy. By adopting human-centered communication strategies, organizations can improve stakeholder engagement, strengthen confidence in AI-assisted recommendations, and promote more inclusive decision-making processes.

From a broader perspective, the proposed communication approach also supports responsible AI adoption by reducing communication barriers between technical experts and organizational decision makers. Enhanced stakeholder understanding contributes to more transparent governance, encourages evidence-based decision making, and minimizes resistance toward AI-supported innovations. These outcomes reflect the humanistic objectives emphasized by Human-Centered AI, where technological advancement is evaluated not only by computational performance but also by its ability to empower individuals, foster collaboration, and improve organizational well-being. The overall relationship among the findings identified in this study is illustrated in Figure 2.

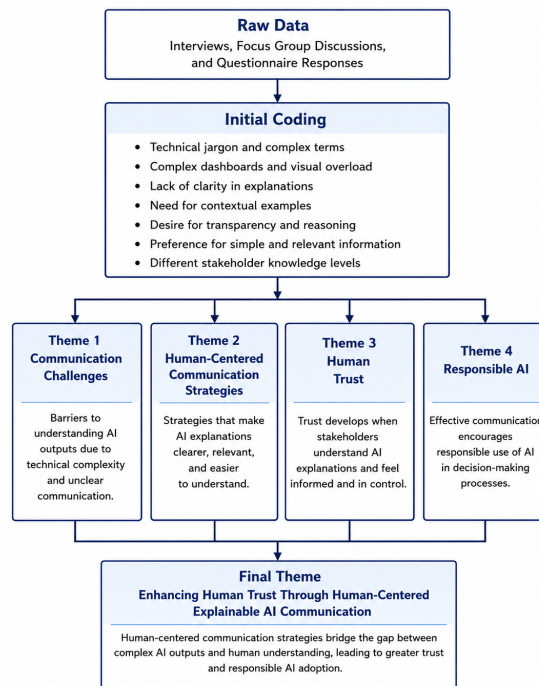


Figure 2. Research Findings Model

Figure 2 summarizes the principal findings of this study by illustrating how explainable AI communication contributes to human trust through a sequential communication process. The model demonstrates that communication effectiveness serves as the initial driver enabling stakeholders to understand AI-generated insights. Improved understanding subsequently strengthens trust, which encourages greater acceptance of AI-assisted recommendations and facilitates responsible organizational decision making. Ultimately, these processes generate broader humanistic organizational outcomes, including improved decision quality, enhanced transparency, inclusive AI adoption, and increased stakeholder confidence. Unlike the conceptual framework presented earlier, which was developed from existing literature, Figure 2 represents the empirical interpretation of the present study and illustrates the practical contribution of human-centered explainable AI communication in organizational contexts.

5. MANAGERIAL IMPLICATIONS

The findings of this study provide practical guidance for organizations seeking to implement XAI in a human-centered manner. Rather than emphasizing algorithmic performance alone, organizations should prioritize communication strategies that enable nontechnical stakeholders to understand AI-generated recommendations with greater clarity and confidence. AI developers and data scientists are encouraged to present

analytical outputs through intuitive visualizations, contextual storytelling, simplified summaries, and audience-centered messaging. These approaches can reduce cognitive complexity, improve stakeholder understanding, and strengthen trust in AI-assisted decision-making across different organizational functions.

From a managerial perspective, organizations should foster collaboration among AI specialists, business managers, and end users to ensure that AI explanations remain transparent, relevant, and aligned with operational needs. Integrating human-centered communication into AI governance can facilitate more inclusive technology adoption, improve decision quality, and enhance organizational transparency. These managerial practices not only support responsible AI implementation but also contribute to broader humanistic outcomes by empowering stakeholders to make informed decisions and promoting sustainable digital transformation consistent with the objectives of SDG 9 (Industry, Innovation and Infrastructure), SDG 16 (Peace, Justice and Strong Institutions), and SDG 17 (Partnerships for the Goals).

6. CONCLUSION

This study explored how explainable AI communication enhances human trust among nontechnical stakeholders through a qualitative investigation of organizational communication practices. The findings demonstrate that trust is not developed solely through the technical performance or transparency of AI systems but is fundamentally influenced by how AI-generated insights are communicated to end users. Four human-centered communication strategies emerged as the principal findings of this study: intuitive data visualization, contextual storytelling, simplified summaries, and audience-centered messaging. Collectively, these strategies improve communication effectiveness, strengthen stakeholder understanding, and facilitate greater trust in AI-assisted decision making.

The study extends the existing Explainable AI literature by positioning communication as a central mechanism linking technical explainability with human trust. While previous studies have primarily emphasized algorithmic transparency and interpretability, the present research demonstrates that communication quality plays an equally important role in enabling stakeholders to confidently understand, evaluate, and utilize AI-generated recommendations. This perspective contributes to the growing body of Human-Centered AI research by integrating communication theory into explainable AI implementation. Despite these contributions, this study has several limitations. The qualitative design limits the generalizability of the findings across different organizational contexts, and participants were selected based on purposive sampling rather than probabilistic techniques. Future research is therefore encouraged to validate the proposed communication framework using quantitative approaches, cross-sector comparisons, or mixed-methods designs involving more diverse stakeholder groups and cultural settings.

Overall, this study concludes that effective explainable AI communication represents a critical enabler of responsible AI adoption. By bridging the gap between complex AI outputs and human understanding, organizations can foster greater stakeholder trust, encourage informed decision making, and generate meaningful humanistic impacts that support the responsible and inclusive development of artificial intelligence.

7. DECLARATIONS

7.1. About Authors

Dwi Apriliasari (DA)  <https://orcid.org/0000-0001-5597-0475>

Bintang Nandana Henry (BN)  <https://orcid.org/0009-0004-7048-6736>

Alexander Williams (AW)  -

7.2. Author Contributions

Conceptualization: DA; Methodology: BN; Software: CH; Validation: MC and AW; Formal Analysis: CH and MC; Investigation: BN; Resources: AW; Data Curation: DA; Writing Original Draft Preparation: BN and CH; Writing Review and Editing: DA and AW; Visualization: MC; All authors, MC, CH, DA, BN, and AW, have read and agreed to the published version of the manuscript.

7.3. Data Availability Statement

The data presented in this study are available on request from the corresponding author.

7.4. Funding

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

7.5. Declaration of Conflicting Interest

The authors declare that they have no conflicts of interest, known competing financial interests, or personal relationships that could have influenced the work reported in this paper.

REFERENCES

- [1] M. A. Rehmat, H. Hassan, M. Rumaan, J. Baig, and K. Abrar, "Human-centred explainable ai in emerging markets: Trust and confidence among non-technical users in pakistan," *Annual Methodological Archive Research Review*, vol. 3, no. 10, pp. 145–163, 2025.
- [2] M. Fahrurrozi, S. R. Djalilah, D. Yunitasari, M. Mohzana, and H. S. Butt, "Development of technopreneurship based e-modules on entrepreneurial learning," *Aptisi Transactions on Technopreneurship (ATT)*, vol. 8, no. 3, pp. 795–808, 2026.
- [3] F. Islam, A. Islam, M. A. Amin, and M. Zaber, "A proposed participatory framework for explainable ai in mhealth: Mixed methods study integrating user and stakeholder requirements," *Journal of Medical Internet Research*, vol. 28, p. e87158, 2026.
- [4] A. Chinnaraju, "Explainable ai (xai) for trustworthy and transparent decision-making: A theoretical framework for ai interpretability," *World Journal of Advanced Engineering Technology and Sciences*, vol. 14, no. 3, pp. 170–207, 2025.
- [5] A. D. Buchdadi, S. Wahyuningsih, Y. Oktavyanti, E. A. Natalia, and H. Zainarthur, "Human centered affective computing models for positive emotional health," *Journal of Orange Technology*, vol. 1, no. 1, pp. 29–38, 2024.
- [6] D. Patil, "Explainable artificial intelligence (xai) for industry applications: Enhancing transparency, trust, and informed decision-making in business operation," *Trust, And Informed Decision-Making In Business Operation (December 03, 2024)*, 2024.
- [7] H. V. Subramanian, C. Canfield, and D. B. Shank, "Designing explainable ai to improve human-ai team performance: a medical stakeholder-driven scoping review," *Artificial Intelligence in Medicine*, vol. 149, p. 102780, 2024.
- [8] A. Rizky, U. Rahardja, M. Muhtarom, D. N. Ramadhan, S. Triandari, and R. F. Terizla, "Responsible ai governance in decentralized web systems," *Legal Autonomous Web-based Governance (LAW)*, vol. 1, no. 1, pp. 123–133, 2026.
- [9] M. Z. Uddin, "Trustworthy ai and explainability," in *Trustworthy Multimodal Intelligent Systems for Independent Living*. Springer, 2025, pp. 111–137.
- [10] J. C. Cheung and S. S. Ho, "The effectiveness of explainable ai on human factors in trust models," *Scientific Reports*, vol. 15, no. 1, p. 23337, 2025.
- [11] A. S. Anita, T. Kuusk, G. Nicola, M. Hardini, and U. Rahardja, "Advancements in artificial intelligence and their contributions to sustainable development goals: A multidisciplinary review," *Smart Human-centered Emerging Research in Machine Intelligence*, vol. 2, no. 1, pp. 37–47, 2026.
- [12] C. Igwe-Nmaju and C. Anadozie, "Commanding digital trust in high-stakes sectors: communication strategies for sustaining stakeholder confidence amid technological risk," *World Journal of Advanced Research and Reviews*, vol. 15, no. 3, pp. 609–630, 2022.
- [13] J. Visave, "Transparency in ai for emergency management: building trust and accountability," *AI and Ethics*, vol. 5, no. 4, pp. 3967–3980, 2025.
- [14] S. L. Sitorus, A. Hermawan, A. W. Widjaja, and M. P. Berlianto, "Enhancing digital fashionpreneurship through innovation capability and market sensing," *Aptisi Transactions on Technopreneurship (ATT)*, vol. 8, no. 2, pp. 645–658, 2026.
- [15] S. Wang, M. A. Qureshi, L. Miralles-Pechuán, T. Huynh-The, T. R. Gadekallu, and M. Liyanage, "Explainable ai for 6g use cases: Technical aspects and research challenges," *IEEE Open Journal of the Communications Society*, vol. 5, pp. 2490–2540, 2024.
- [16] M. M. Sari, U. Rahardja, N. Azizah, G. Fransiso, and J. Bagaskara, "Analysis of academic information system integration in improving the quality of higher education data reporting to pddikti," *International Transactions on Education Technology (ITEE)*, vol. 4, no. 2, pp. 172–186, 2026.

- [17] F. Islam, A. Islam, M. A. Amin, and M. Zaber, "A proposed participatory framework for explainable ai in mhealth: Mixed methods study integrating user and stakeholder requirements," *Journal of Medical Internet Research*, vol. 28, p. e87158, 2026. [Online]. Available: <https://pmc.ncbi.nlm.nih.gov/articles/PMC13138816/>
- [18] N. Polemi, I. Praça, K. Kioskli, and A. Bécue, "Challenges and efforts in managing ai trustworthiness risks: a state of knowledge," *Frontiers in big Data*, vol. 7, p. 1381163, 2024.
- [19] Z. Atf and P. R. Lewis, "Human centricity in the relationship between explainability and trust in ai," *IEEE Technology and Society Magazine*, vol. 42, no. 4, pp. 66–76, 2024.
- [20] M. A. K. Akhtar, M. Kumar, and A. Nayyar, "The role of human-centered design in developing explainable ai," in *Towards Ethical and Socially Responsible Explainable AI: Challenges and Opportunities*. Springer, 2024, pp. 99–126.
- [21] H. Shehu, E. Ogunleye, M. O. Atilola, E. I. Eromosele, A. B. Lawal, and T. T. Chukwuma, "Ethical and responsible ai in engineering and construction projects: Governance, trust, and human-centered design," *Scientific Journal of Engineering, and Technology*, vol. 2, no. 2, pp. 53–62, 2025.
- [22] M. Szymanski, V. Vanden Abeele, and K. Verbert, "Disentangling stakeholder role and expertise in user-centered explainable ai," in *Proceedings of the 33rd ACM Conference on User Modeling, Adaptation and Personalization*, 2025, pp. 32–39.
- [23] H. Sun, Y. Liu, A. Al-Tahmeesschi, A. Nag, M. Soleimanpour, B. Canberk, H. Arslan, and H. Ahmadi, "Advancing 6g: Survey for explainable ai on communications and network slicing," *IEEE Open Journal of the Communications Society*, vol. 6, pp. 1372–1412, 2025.
- [24] B. Teixeira, L. Carvalhais, T. Pinto, and Z. Vale, "Explainable ai framework for reliable and transparent automated energy management in buildings," *Energy and Buildings*, p. 116246, 2025.
- [25] A. Jayaloka and A. Yohannis, "Enhancing business intelligence with explainable ai: Evaluating transparency, interpretability, and user trust," *Buletin Poltanesa*, vol. 25, no. 1, 2025.
- [26] T. L. Anita, R. G. Freedy, N. Silawati, and M. Sunengsih, "Mathematical logic as a foundation for ai-driven decision-making systems," *Smart Human-centered Emerging Research in Machine Intelligence*, vol. 2, no. 1, pp. 14–25, 2026.
- [27] W. Xu, "Human-centered artificial intelligence (hcai): Foundations and approaches," *arXiv preprint arXiv:2601.01247*, 2026.
- [28] M. Camacho, J. Perramon, X. Puig-Bosch, V. N. Dang, O. Díaz, and K. Lekadir, "Stakeholder engagement: the path to trustworthy ai in healthcare," in *Trustworthy AI in Medical Imaging*. Elsevier, 2025, pp. 471–493.
- [29] F. Herrera, "Intelligent organizational engineering driven by human-ai collaboration and explainable ai to increase productivity," *Dirección y Organización*, no. 87, pp. 5–14, 2025.
- [30] E. Arif, S. Suherman, and A. P. Widodo, "Analyzing public sentiment on digital banks in indonesia via social media x," *Aptisi Transactions on Technopreneurship (ATT)*, vol. 8, no. 1, pp. 253–267, 2026.
- [31] A. Mestre, M. Albert, M. Gil, and V. Pelechano, "Explainability across the spectrum: Modeling stakeholder goals based on ai complexity levels," in *2025 IEEE 33rd International Requirements Engineering Conference (RE)*. IEEE, 2025, pp. 544–551.
- [32] P. K. Myakala and A. K. Jonnalagadda, "Explainable ai for sustainability: Bridging trust, ethics, and accountability," in *Recent Advances in Artificial Intelligence for Sustainable Development (RAISD 2025)*. Atlantis Press, 2025, pp. 563–576.
- [33] M. D. M. Qureshi, M. A. Qureshi, and W. Rashwan, "Explainable ai for hate speech moderation: A stakeholder-centered and sociotechnical review," *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, vol. 16, no. 1, p. e70076, 2026.
- [34] S. Bobek, P. Korycińska, M. Krakowska, M. Mozolewski, D. Rak, M. Zych, M. Wójcik, and G. J. Nalepa, "User-centric evaluation of explainability of ai with and for humans: a comprehensive empirical study," *International Journal of Human-Computer Studies*, p. 103625, 2025.
- [35] O. Emi-Johnson, O. Fasanya, and A. Adeniyi, "Explainable ai for pesticide decision-making: Enhancing trust in datadriven crop protection models," *International Journal of Computer Applications Technology and Research*, vol. 14, no. 01, pp. 147–161, 2025.
- [36] M. Akintunde, V. Young, V. Yazdanpanah, A. Salehi Fathabadi, P. Leonard, M. Butler, and L. Moreau, "Verifiably safe and trusted human-ai systems: a socio-technical perspective," in *Proceedings of the First International Symposium on Trustworthy Autonomous Systems*, 2023, pp. 1–6.

- [37] C. Agostinho, Z. Dikopoulou, E. Lavasa, K. Perakis, S. Pitsios, R. Branco, S. Reji, J. Hetterich, E. Biliri, F. Lampathaki *et al.*, “Explainability as the key ingredient for ai adoption in industry 5.0 settings,” *Frontiers in artificial intelligence*, vol. 6, p. 1264372, 2023.
- [38] E. Nyong, “Explainable ai (xai) for high-stakes decision systems,” *International Journal of Technology, Management and Humanities*, vol. 11, no. 02, pp. 118–129, 2025.
- [39] F. Herrera and R. Calderón, “Opacity as a feature, not a flaw: Role-sensitive explainability, institutional trust, and the lobox ethics governance framework for ai,” *Technology in Society*, p. 103302, 2026.
- [40] N. Azizah, P. A. Sunarya, U. Rahardja, A. B. Mutiara, P. Prihandoko, and C. Pasha, “Improving smear-negative tuberculosis detection using data augmentation and faster r-cnn,” *International Journal of Cyber and IT Service Management (IJCITSM)*, vol. 6, no. 1, pp. 65–77, 2026.
- [41] M. Kinney, M. Anastasiadou, M. Naranjo-Zolotov, and V. Santos, “Expectation management in ai: A framework for understanding stakeholder trust and acceptance of artificial intelligence systems,” *Heliyon*, vol. 10, no. 7, 2024.
- [42] O. Olateju, S. U. Okon, O. O. Olaniyi, A. D. Samuel-Okon, and C. U. Asonze, “Exploring the concept of explainable ai and developing information governance standards for enhancing trust and transparency in handling customer data,” *Available at SSRN*, 2024.
- [43] R. Vedashree, J. Sahiba, and B. Agarwal, “Towards trustworthy ai: Guidelines for operationalization and responsible,” *The Quest for AI Sovereignty, Transparency and Accountability: Official Outcome of the UN IGF Data and Artificial Intelligence Governance Coalition*, p. 89, 2026.
- [44] A. Chaudhary, N. L. Rane, R. Rastogi, and J. Rane, “Explainable and ethical artificial intelligence in customer relationship management: Improving customer experience, value creation, and loyalty,” *International Journal of Applied Resilience and Sustainability*, vol. 2, no. 2, pp. 243–278, 2026.
- [45] N. I. Susanthi, M. Ali, and A. H. Hernawan, “Digital learning platforms as facilitator for university-business collaboration in logistics management curriculum design,” *International Journal of Cyber and IT Service Management (IJCITSM)*, vol. 6, no. 1, pp. 37–50, 2026.
- [46] N. I. Susanthi, M. F. Djamaly, A. Fitriani, M. Mardiana, and C. T. Hua, “Design and evaluation of emotionally adaptive chatbots to promote positive mental well-being in young adults,” *Journal of Orange Technology*, vol. 1, no. 2, pp. 51–62, 2025.
- [47] M. Fairhurst, K. Kaesling, and V. Klös, “Legal requirements, trust issues and engineering challenges-a multi-disciplinary case for user-specific explainability,” in *World Conference on Explainable Artificial Intelligence*. Springer, 2025, pp. 354–377.
- [48] A. D. Sunny, “Trust in transparency: How explainable ai shapes user perceptions,” *arXiv preprint arXiv:2510.04968*, 2025.
- [49] A. Taneja, “Explainable ai: Building trust and transparency in machine learning models,” *Journal Of Engineering And Computer Sciences*, vol. 4, no. 8, pp. 482–489, 2025.
- [50] M. H. Lee, S. X. Y. Choo, S. D. Thilarajah *et al.*, “Improving health professionals’ onboarding with ai and xai for trustworthy human-ai collaborative decision making,” *arXiv preprint arXiv:2405.16424*, 2024.
- [51] M. F. Hilmi, H. Hamdan, A. Faturahman, B. N. Henry, and J. Wilson, “Ai in industry forecasting: The use of ai in predicting industry trends and demands,” *Health, Empathy, and AI Learning (HEAL)*, vol. 1, no. 1, pp. 78–83, 2025.
- [52] P. Tsarouhas and K. Grigoriadis, “Building trust in ai for public administration: A strategic framework for transparency, xai, participation, and digital literacy,” in *2025 7th International Congress on Human-Computer Interaction, Optimization and Robotic Applications (ICHORA)*. IEEE, 2025, pp. 1–9.
- [53] A. Nastoska, B. Jancheska, M. Rizinski, and D. Trajanov, “Evaluating trustworthiness in ai: Risks, metrics, and applications across industries,” *Electronics*, vol. 14, no. 13, p. 2717, 2025.