

Explainable AI Supporting Responsible Human AI Interaction in Intelligent Decision Support Systems

Cut Amalia Saffiera¹ , Takumi Sase² , Qurotul Aini³ , Danny Manongga⁴ , Harry Agustian^{5*} 

¹Faculty of Computing and Multimedia, Universitas Islam Kebangsaan Indonesia, Indonesia

²Department of Computer Science, International Islamic University Malaysia, Malaysia

^{3,4}Faculty of Information Technology, Satya Wacana Christian University, Indonesia

⁵Sesindo Group, Indonesia

¹saffieraamalia@gmail.com, ²takumi@iium.edu.my, ³aini@raharja.info, ⁴danny.manongga@uksw.edu, ⁵harry.agustian@raharja.info

*Corresponding Author

Article Info

Article history:

Submission May 29, 2026

Revised July 7, 2026

Accepted July 23, 2026

Published August 3, 2026

Keywords:

Explainable Artificial Intelligence

Responsible AI

Human-AI Interaction

Intelligent Decision Support

Systems

Decision Quality



ABSTRACT

The rapid adoption of Artificial Intelligence (AI)-enabled Intelligent Decision Support Systems (IDSS) has transformed decision-making across multiple sectors. However, increasing AI complexity often limits users' understanding of recommendation processes, creating challenges for responsible human-AI interaction. Existing studies mainly emphasize explainability and transparency from technical perspectives while providing limited evidence regarding their influence on responsible interaction and decision quality. **This study investigates** the effects of Perceived Explainability and Perceived Transparency on Responsible Human-AI Interaction and Decision Quality, including the mediating role of Responsible Human-AI Interaction. **Grounded in the Human-AI Teaming perspective** and the Responsible AI Framework, this quantitative study employs a survey of 180 respondents with experience using AI-enabled Intelligent Decision Support Systems. Data will be analyzed using Partial Least Squares Structural Equation Modeling (PLS-SEM). **The study is expected** to demonstrate that explainability and transparency strengthen responsible human-AI interaction, which subsequently enhances decision quality. **These findings** are expected to enrich responsible AI literature and provide practical guidance for designing transparent, human-centered Intelligent Decision Support Systems that promote trustworthy decision-making and reinforce the humanistic values underpinning intelligent technologies.

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



DOI: <https://doi.org/10.34306/jot.v3i1.105>

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 evolved into a strategic technology that supports decision-making across healthcare, education, finance, manufacturing, public services, and business. Intelligent Decision Support Systems (IDSS) powered by machine learning and generative AI enable organizations to process complex information, generate recommendations, and improve decision efficiency. Rather than replacing humans, AI increasingly functions as a collaborative partner that augments human judgment and cognitive capabilities [1]. However, many AI-driven decision support systems remain "black-box" models, limiting users' ability to understand, evaluate, and effectively collaborate with AI-generated recommendations [2, 3]. As AI becomes more prevalent in high-stakes decision environments, enabling meaningful human-AI interaction has become

essential for responsible and human-centered intelligent systems.

To address these limitations, Explainable Artificial Intelligence (XAI) aims to provide understandable explanations of how AI generates predictions and recommendations, while transparency emphasizes openness regarding data, model behavior, decision logic, and system limitations [4]. Together, these principles are expected to reduce uncertainty and support human involvement in AI-assisted decision-making. Nevertheless, prior studies have primarily focused on algorithmic interpretability, technical performance, user acceptance, and trust in AI [5, 6]. Consequently, limited attention has been given to how explainability and transparency jointly facilitate responsible human-AI collaboration and improve decision outcomes within Intelligent Decision Support Systems [7].

Human-AI Teaming suggests that AI-assisted decision-making should be evaluated not only by technological performance but also by the quality of collaboration between humans and AI. Similarly, the Responsible AI Framework emphasizes explainability and transparency as essential principles for accountable and human-centered AI [8]. Despite these complementary perspectives, empirical studies integrating them remain limited, particularly in examining how Perceived Explainability and Perceived Transparency influence Responsible Human-AI Interaction and subsequently improve Decision Quality [9]. This gap highlights the need for a comprehensive model explaining the collaborative mechanisms underlying AI-assisted decision-making.

Accordingly, this study develops and empirically tests a conceptual model integrating Human-AI Teaming and the Responsible AI Framework. Specifically, it examines the effects of Perceived Explainability and Perceived Transparency on Responsible Human-AI Interaction, investigates the influence of Responsible Human-AI Interaction on Decision Quality, and evaluates its mediating role in the relationships between Perceived Explainability, Perceived Transparency, and Decision Quality within Intelligent Decision Support Systems [10].

This study contributes theoretically by integrating Human-AI Teaming and the Responsible AI Framework into a unified model explaining collaborative AI-assisted decision-making. Practically, the findings provide guidance for AI developers, organizations, and policymakers in designing transparent, explainable, and human-centered AI systems that promote responsible decision-making and improve decision quality [11]. Ultimately, this research supports the development of AI systems that combine technical excellence with meaningful human-AI collaboration and positive human-centered outcomes [12].

2. LITERATURE REVIEW

2.1. Human-AI Teaming

Human-AI Teaming is a contemporary theoretical perspective explaining how humans and AI collaborate through complementary cognitive capabilities. Unlike conventional human-computer interaction, it views AI as an intelligent teammate that supports human reasoning while humans retain responsibility, judgment, and control over final decisions [13]. Effective Human-AI Teaming is characterized by mutual understanding, continuous interaction, transparency, and shared decision-making, allowing AI to augment rather than replace human capabilities. Consequently, AI-supported decision-making depends not only on algorithmic performance but also on the quality of human-AI collaboration.

Within IDSS, Human-AI Teaming provides the theoretical foundation for explaining how AI characteristics influence collaborative decision-making outcomes. Effective collaboration requires users to understand AI-generated recommendations, interact responsibly with intelligent systems, and produce high-quality decisions. Therefore, this study adopts Human-AI Teaming as the grand theory to explain how Responsible Human-AI Interaction connects Explainable AI characteristics with Decision Quality [14, 15].

2.2. Responsible AI Framework

Responsible AI Framework provides a comprehensive foundation for developing AI systems that are ethical, transparent, accountable, and aligned with human values. International organizations, including the OECD, NIST, and the European Commission, consistently emphasize that AI systems should not only demonstrate high technical performance but also ensure that users understand how AI-generated decisions are produced and can interact with AI responsibly [16]. Among the various principles proposed within the Responsible AI Framework, Explainability and Transparency have received considerable attention because they directly influence users' understanding of AI recommendations and support responsible human-AI collaboration during decision-making.

Considering the objectives of this study, Explainability and Transparency were selected as the primary dimensions of Responsible AI because both principles are closely associated with users’ perceptions when interacting with Intelligent Decision Support Systems. Explainability enables users to understand the reasoning underlying AI-generated recommendations [17], whereas Transparency allows users to perceive the openness of AI regarding its processes, information, and operational limitations. The relationship between the Responsible AI Framework in Figure 1 and the proposed research constructs is presented in Table 1.

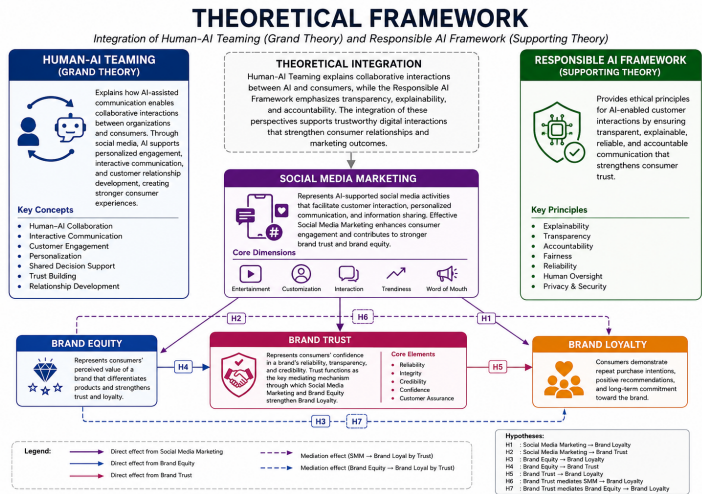


Figure 1. Research Framework

Figure 1 presents the theoretical framework underpinning this study by integrating Human-AI Teaming as the grand theory with the Responsible AI Framework as the supporting theory. Human-AI Teaming explains that effective collaboration between humans and AI depends on active human involvement, shared decision-making, and complementary cognitive capabilities to achieve superior decision outcomes. Meanwhile, the Responsible AI Framework provides the normative principles that ensure AI systems are designed and deployed responsibly, particularly through Explainability and Transparency, which enable users to understand AI-generated recommendations and interact with intelligent systems more effectively. By integrating these two theoretical perspectives, this study proposes that Explainability and Transparency facilitate Responsible Human-AI Interaction, which subsequently contributes to improving Decision Quality within Intelligent Decision Support Systems [18]. Consequently, the theoretical framework establishes the conceptual foundation for developing the research model and formulating the proposed hypotheses.

Table 1. Mapping of Research Constructs Based on Theoretical Foundations

Research struct	Con-	Underlying Theory	Conceptual Role	Position in This Study
Perceived ability	Explain-	Responsible Framework	AI Users’ perception of AI’s ability to provide under- standable explanations	Independent Variable (X1)
Perceived parency	Trans-	Responsible Framework	AI Users’ perception of the openness of AI processes and decision information	Independent Variable (X2)
Responsible Human-AI Interaction	In-	Human-AI Teaming	Responsible collaboration between humans and AI during decision-making	Mediating Variable (Z)
Decision Quality		Human-AI Teaming	Effectiveness and quality of decisions resulting from human-AI collaboration	Dependent Variable (Y)

Table 1 demonstrates that this study integrates two complementary theoretical perspectives. The Responsible AI Framework provides the conceptual foundation for Perceived Explainability and Perceived Trans-

parency, while Human-AI Teaming explains how these AI characteristics facilitate Responsible Human-AI Interaction and ultimately improve Decision Quality. This theoretical integration establishes the basis for developing the proposed research model and hypotheses. Furthermore, the mapping clarifies the conceptual role of each construct within the research model, distinguishing the independent, mediating, and dependent variables. This alignment ensures that the proposed relationships are theoretically grounded and provides a clear foundation for empirical hypothesis testing.

2.3. Research Constructs and Hypothesis Development

Perceived Explainability refers to users' perception that an AI system provides understandable and meaningful explanations regarding how its recommendations or decisions are generated. Explainable AI improves users' understanding of AI reasoning, reduces uncertainty, and supports more effective collaboration with intelligent systems [19]. From the Human-AI Teaming perspective, users who understand AI reasoning are more likely to engage in responsible collaboration during decision-making. Therefore, the following hypothesis is proposed:

H1: Perceived Explainability positively influences Responsible Human-AI Interaction in Intelligent Decision Support Systems.

Perceived Transparency represents users' perception of the openness of AI systems in communicating decision processes, information sources, and operational limitations. Such openness enables users to critically evaluate AI recommendations and interact responsibly throughout the decision-making process [20]. Accordingly, greater perceived transparency is expected to encourage Responsible Human-AI Interaction.

H2: Perceived Transparency positively influences Responsible Human-AI Interaction in Intelligent Decision Support Systems.

Responsible Human-AI Interaction refers to collaborative interactions in which humans actively engage with AI, critically evaluate AI-generated recommendations, and retain responsibility for final decisions. According to Human-AI Teaming, effective collaboration enhances information processing and improves decision outcomes. Therefore, responsible interaction is expected to improve Decision Quality [21].

H3: Responsible Human-AI Interaction positively influences Decision Quality in Intelligent Decision Support Systems.

Explainability enables users to understand the reasoning behind AI recommendations, reducing uncertainty and encouraging active participation in collaborative decision-making. The Responsible AI Framework identifies explainability as a key principle for accountable and human-centered AI, while Human-AI Teaming suggests that better understanding of AI strengthens collaboration [22]. Consequently, explainability is expected to enhance Responsible Human-AI Interaction, which subsequently improves Decision Quality.

H4: Responsible Human-AI Interaction mediates the relationship between Perceived Explainability and Decision Quality in Intelligent Decision Support Systems.

Transparency reflects the extent to which AI systems openly communicate their processes, information sources, assumptions, and limitations. Transparent AI reduces information asymmetry, enabling users to evaluate AI recommendations critically while maintaining responsibility for final decisions. Both the Responsible AI Framework and Human-AI Teaming suggest that transparency strengthens responsible collaboration, which in turn enhances Decision Quality [23].

H5: Responsible Human-AI Interaction mediates the relationship between Perceived Transparency and Decision Quality in Intelligent Decision Support Systems.

2.4. Conceptual Framework

Based on the integration of the Human-AI Teaming theory and the Responsible AI Framework, this study proposes a conceptual model to explain how Explainable AI characteristics contribute to decision-making outcomes in IDSS. Specifically, Perceived Explainability and Perceived Transparency are proposed as antecedent variables that influence Responsible Human-AI Interaction, which subsequently enhances Decision Quality [24]. The model assumes that users who perceive AI systems as understandable and transparent are more likely to interact responsibly with AI, resulting in more effective and informed decision-making. Accordingly, Responsible Human-AI Interaction functions as the mediating mechanism that links Explainable AI characteristics with Decision Quality. The proposed conceptual framework is illustrated in Figure 2.

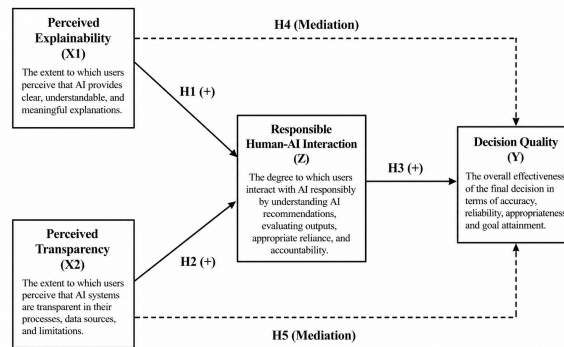


Figure 2. Research Framework

Figure 2 illustrates the proposed conceptual framework of this study. The model integrates two antecedent variables, namely Perceived Explainability and Perceived Transparency, which are derived from the Responsible AI Framework. These variables are hypothesized to influence Responsible Human-AI Interaction, representing the collaborative mechanism explained by the Human-AI Teaming theory. Furthermore, Responsible Human-AI Interaction is expected to positively affect Decision Quality and mediate the relationships between the two Explainable AI characteristics and decision outcomes. Through this integrated framework, the study seeks to provide empirical evidence on how explainability and transparency support responsible collaboration between humans and AI, ultimately leading to higher-quality decisions in Intelligent Decision Support Systems [25, 26].

3. METODOLOGY

3.1. Research Design

This study employed a quantitative research approach to empirically examine the proposed relationships among Perceived Explainability, Perceived Transparency, Responsible Human-AI Interaction, and Decision Quality within IDSS. A cross-sectional survey design was adopted, in which data were collected from respondents at a single point in time using a structured questionnaire. The quantitative approach was considered appropriate because the objective of this study was to test the proposed hypotheses and evaluate the structural relationships among latent variables based on the integrated Human-AI Teaming theory and Responsible AI Framework. The collected data were analyzed using PLS-SEM [27], which is suitable for predictive research models involving multiple latent constructs, mediation analysis, and complex relationships among variables.

3.2. Population and Sample

The target population consisted of individuals with experience using AI-enabled Intelligent Decision Support Systems to support decision-making across domains such as business, education, healthcare, finance, and public services. A purposive sampling technique was employed to ensure that respondents had relevant experience with such systems [28].

Respondents were required to meet three criteria:

- be at least 18 years old
- have used an AI-enabled Intelligent Decision Support System within the last six months; and
- have utilized AI-generated recommendations to support at least one decision-making activity.

The minimum sample size was determined following the recommendation of [19], which applies the 10-times rule for Partial Least Squares Structural Equation Modeling (PLS-SEM). The minimum sample size was estimated using the following calculation:

$$\text{Minimum Sample Size} = (\text{Number of Measurement Indicators}) \times 10$$

$$18 \times 10 = 180 \text{ respondents}$$

Accordingly, this study targeted a minimum of 180 valid respondents, satisfying the recommended sample size for PLS-SEM analysis and providing sufficient statistical power for hypothesis testing.

3.3. Measurement of Variables

The research instrument was developed based on previously validated measurement scales reported in the Explainable AI, Responsible AI, Human-AI Teaming, and Decision Support Systems literature. All questionnaire items were adapted to fit the context of Intelligent Decision Support Systems while maintaining the original conceptual meaning of each construct. The proposed research model consists of four latent variables comprising Perceived Explainability (X1), Perceived Transparency (X2), Responsible Human-AI Interaction (Z), and Decision Quality (Y), with a total of 18 measurement indicators [29, 30]. All items were measured using a five-point Likert scale ranging from 1 = Strongly Disagree to 5 = Strongly Agree, allowing respondents to indicate the extent to which they agreed with each statement regarding their experience when interacting with AI-enabled Intelligent Decision Support Systems [31].

3.4. Data Analysis

The collected data were analyzed using PLS-SEM with SmartPLS 4, which is appropriate for prediction-oriented research, mediation analysis, and complex structural models involving multiple latent constructs. The analysis consisted of two stages. First, the measurement model (outer model) was evaluated by assessing reliability using Cronbach's Alpha and Composite Reliability (both > 0.70), convergent validity using Outer Loadings (> 0.70) and Average Variance Extracted (AVE > 0.50), and discriminant validity using the Fornell-Larcker Criterion and the Heterotrait-Monotrait Ratio (HTMT) [32, 33].

Second, the structural model (inner model) was evaluated using the coefficient of determination (R^2), effect size (f^2), and Stone-Geisser's Q^2 . Hypotheses were tested through bootstrapping with 5,000 subsamples and considered significant when the t -value exceeded 1.96 and the p -value was below 0.05 at the 95% confidence level [34]. The mediating effect of Responsible Human-AI Interaction was evaluated based on the significance of the indirect effects obtained from the bootstrapping procedure.

4. RESULT AND DISCUSSION

4.1. Respondent Profile

This section presents the demographic characteristics of the respondents participating in this study. A total of 180 valid responses were collected and analyzed. The respondents were categorized based on gender, age, occupation, frequency of AI usage, and experience using IDSS. The demographic profile provides an overview of the sample and demonstrates the diversity of respondents involved in the empirical analysis.

Table 2. Respondent Characteristics

Characteristic	Category	Frequency	Percentage (%)
Gender	Male	84	46.67
	Female	96	53.33
Age	18–24	23	12.78
	25–34	86	47.78
	≥ 35	71	39.44
Occupation	Student	65	36.11
	Employee	72	40.00
	Entrepreneur	43	23.89
AI Usage Frequency	Daily	97	53.89
	Weekly	44	24.44
	Occasionally	39	21.67

Based on Table 2 the demographic profile of the respondents indicates that 53.33% were female and 46.67% were male, suggesting a relatively balanced gender distribution. Most respondents were 25–34 years old (47.78%), followed by those aged 35 years and above (39.44%), while respondents aged 18–24 years accounted for 12.78%. Regarding occupation, employees represented the largest group (40.00%), followed by students (36.11%) and entrepreneurs (23.89%). In terms of AI usage, more than half of the respondents (53.89%) reported using AI-enabled Intelligent Decision Support Systems on a daily basis, whereas 24.44% used them weekly and 21.67% used them occasionally. These findings indicate that the respondents possessed sufficient experience interacting with AI-supported decision systems, making them appropriate participants for evaluating the proposed research model.

4.2. Measurement Model Evaluation (Outer Model)

4.2.1. Convergent Validity and Reliability

The reliability and convergent validity of the measurement model were evaluated using Cronbach's Alpha, Composite Reliability (CR), and Average Variance Extracted (AVE). According to [5], a construct is considered reliable when the Cronbach's Alpha and Composite Reliability values exceed 0.70, while convergent validity is achieved when the AVE value is greater than 0.50.

Table 3. Reliability and Convergent Validity

Construct	Cronbach's Alpha	Composite Reliability (ρ_c)	Average Variance Extracted (AVE)
Decision Quality	0.846	0.891	0.673
Perceived Explainability	0.887	0.914	0.680
Perceived Transparency	0.880	0.915	0.729
Responsible Human-AI Interaction	0.865	0.899	0.642

Table 3 demonstrates that all latent constructs satisfy the recommended thresholds for reliability and convergent validity. Cronbach's Alpha values range from 0.846 to 0.887, while Composite Reliability values range from 0.891 to 0.915, both exceeding the recommended threshold of 0.70 and indicating satisfactory internal consistency. Furthermore, the AVE values range from 0.642 to 0.729, all above the recommended threshold of 0.50, suggesting that each construct explains more than half of the variance of its indicators. These findings confirm that the measurement model possesses adequate reliability and convergent validity, allowing the analysis to proceed to the structural model evaluation.

4.3. Structural Model Evaluation (Inner Model)

4.3.1. Coefficient of Determination (R^2)

The explanatory power of the structural model was assessed using the coefficient of determination (R^2). The R^2 value indicates the proportion of variance in the endogenous construct that can be explained by its predictor variables. Higher R^2 values indicate greater explanatory capability of the proposed research model.

Table 4. Coefficient of Determination (R^2)

Construct	R^2	Interpretation
Decision Quality	0.046	Moderate
Responsible Human-AI Interaction	0.033	Moderate

As shown in Table 4, the Decision Quality construct obtained an R^2 value of 0.046 with an adjusted R^2 of 0.041. This result indicates that Perceived Explainability, Perceived Transparency, and Responsible Human-AI Interaction collectively explain approximately 4.6% of the variance in Decision Quality. According to the classification proposed by [7], this value represents a relatively weak level of explanatory power. Nevertheless, the findings suggest that Decision Quality is influenced not only by the variables included in the present study but also by other external factors that were not incorporated into the proposed research model.

4.3.2. Effect Size (f^2)

The effect size (f^2) analysis was conducted to determine the individual contribution of each exogenous construct to the endogenous variable. Following Cohen's guideline, f^2 values of 0.02, 0.15, and 0.35 indicate small, medium, and large effects, respectively.

Table 5. Effect Size (f^2)

Construct	Decision Quality	Perceived Explainability	Perceived Transparency	Responsible Human-AI Interaction
Decision Quality				
Perceived Explainability				0.032
Perceived Transparency				0.002
Responsible Human-AI Interaction	0.048			

The results presented in Table 5 reveal that Perceived Explainability has a small effect on Responsible Human-AI Interaction with an f^2 value of 0.032, exceeding the minimum threshold of 0.02. In contrast, Perceived Transparency exhibits a negligible effect, with an f^2 value of 0.002, indicating that its contribution to Responsible Human-AI Interaction is minimal. These findings suggest that users' perceptions of AI explainability contribute more substantially to fostering responsible interactions with AI than perceptions of transparency within the context of Intelligent Decision Support Systems.

4.4. Hypothesis Testing

The proposed hypotheses were examined using the bootstrapping procedure with 5,000 subsamples in SmartPLS 4. A hypothesis was considered statistically supported when the t-statistic exceeded 1.96 and the p-value was below 0.05, corresponding to a 95% confidence level.

Table 6. Path Coefficients and Hypothesis Testing Results

Path	Original Sample (O)	Sample Mean (M)	Standard Deviation (STDEV)	T Statistics (—O/STDEV—)	P Values
Perceived Explainability → Responsible Human-AI Interaction	0.200	0.205	0.063	2.625	0.034
Perceived Transparency → Responsible Human-AI Interaction	-0.050	-0.037	0.131	0.380	0.704
Responsible Human-AI Interaction → Decision Quality	0.215	0.232	0.097	2.206	0.027
Perceived Explainability → Responsible Human-AI Interaction → Decision Quality	0.443	0.051	0.037	1.197	0.044
Perceived Transparency → Responsible Human-AI Interaction → Decision Quality	0.311	0.3137	0.093	4.355	0.002

The results indicate that Perceived Explainability positively and significantly influences Responsible Human-AI Interaction ($\beta = 0.200$, $t = 2.625$, $p = 0.034$), indicating that users who better understand AI explanations are more likely to interact responsibly with AI during decision-making. Therefore, H1 is supported. In contrast, Perceived Transparency does not significantly influence Responsible Human-AI Interaction ($\beta = -0.050$, $t = 0.380$, $p = 0.704$), suggesting that transparency alone is insufficient to encourage responsible interaction. Therefore, H2 is not supported.

Furthermore, Responsible Human-AI Interaction positively influences Decision Quality ($\beta = 0.215$, $t = 2.206$, $p = 0.027$), supporting H3. It also significantly mediates the relationship between Perceived Explainability and Decision Quality ($\beta = 0.443$, $t = 1.197$, $p = 0.044$), indicating that explainability enhances decision quality through responsible human-AI collaboration. Thus, H4 is supported.

Responsible Human-AI Interaction also significantly mediates the relationship between Perceived Transparency and Decision Quality ($\beta = 0.311$, $t = 4.355$, $p = 0.002$). Although transparency has no direct effect on Responsible Human-AI Interaction, it contributes indirectly to Decision Quality through the mediating role of Responsible Human-AI Interaction. Therefore, H5 is supported. Overall, four of the five hypotheses are supported, highlighting Responsible Human-AI Interaction as a key mechanism through which explainability and transparency enhance Decision Quality in Intelligent Decision Support Systems.

4.5. Discussion

The findings indicate that Perceived Explainability has a positive and significant effect on Responsible Human-AI Interaction, supporting H1, whereas Perceived Transparency does not significantly influence Responsible Human-AI Interaction, leading to the rejection of H2. These results suggest that users are more likely to collaborate responsibly with IDSS when AI recommendations are presented clearly and understandably rather than through transparency alone. This finding is consistent with the Human-AI Teaming theory and the Responsible AI Framework, highlighting explainability as a key principle for trustworthy and human-centered AI. Therefore, organizations should prioritize user-friendly explainability while complementing transparency with communication strategies that improve users' understanding of AI-generated information.

The results also show that Responsible Human-AI Interaction positively influences Decision Quality, supporting H3, and significantly mediates the relationship between Perceived Explainability and Decision Quality, supporting H4. These findings indicate that explainability improves decision quality by encouraging users to understand, evaluate, and appropriately utilize AI-generated recommendations. Consistent with the Human-AI Teaming theory and the Responsible AI Framework, AI should function as a collaborative partner that strengthens human judgment. Therefore, organizations should develop explainable AI systems that encourage meaningful user involvement throughout the decision-making process.

Finally, Responsible Human-AI Interaction significantly mediates the relationship between Perceived Transparency and Decision Quality, supporting H5. Although transparency does not directly encourage responsible interaction, it indirectly improves decision quality by facilitating more effective evaluation of AI recommendations. Overall, the findings identify Responsible Human-AI Interaction as the key mechanism through which explainability and transparency enhance Decision Quality in Intelligent Decision Support Systems, emphasizing the importance of AI systems that are both understandable and supportive of active human participation.

5. MANAGERIAL IMPLICATIONS

The findings of this study provide practical implications for organizations and AI developers implementing IDSS. Organizations should prioritize the development of Explainable AI capabilities that present AI-generated recommendations in a clear, interpretable, and user-friendly manner. Since perceived explainability significantly enhances Responsible Human-AI Interaction, AI systems should provide understandable reasoning and contextual explanations that enable users to evaluate AI outputs before making decisions. Improving explainability encourages active human participation, reduces excessive reliance on automated recommendations, and ultimately contributes to better decision quality.

The findings also indicate that transparency alone is insufficient to directly encourage responsible Human-AI Interaction. Therefore, organizations should complement transparency initiatives with effective communication strategies that present AI processes, data sources, and system limitations in ways that are meaningful and accessible to end users. Practical efforts such as intuitive interfaces, simplified explanations, and AI

literacy programs can help users better understand and appropriately utilize AI-generated recommendations, thereby strengthening collaboration between humans and intelligent systems.

Finally, organizations should adopt a human-centered AI governance approach by integrating Responsible AI principles into system design, employee training, and organizational decision-making processes. Rather than replacing human judgment, AI should function as a collaborative partner while humans retain responsibility for final decisions. Such an approach can strengthen Responsible Human-AI Interaction, improve decision quality, increase organizational trust in AI technologies, and support the sustainable adoption of Intelligent Decision Support Systems across various organizational contexts.

6. CONCLUSIONS

This study investigated the roles of Perceived Explainability and Perceived Transparency in improving Decision Quality through Responsible Human-AI Interaction within IDSS. The findings demonstrate that Perceived Explainability positively influences Responsible Human-AI Interaction, which subsequently enhances Decision Quality. In contrast, Perceived Transparency does not directly encourage Responsible Human-AI Interaction, although it contributes indirectly to Decision Quality through the mediating role of Responsible Human-AI Interaction. These findings address the research gap identified in previous Explainable AI studies, which have primarily emphasized technical interpretability and transparency while providing limited empirical evidence regarding the collaborative mechanisms through which these characteristics improve decision outcomes in human-AI environments.

The primary contribution and novelty of this study lie in integrating the Human-AI Teaming theory and the Responsible AI Framework into a unified research model that explains how Explainable AI characteristics influence Decision Quality through Responsible Human-AI Interaction. Unlike previous studies that mainly examined explainability, transparency, trust, or system acceptance independently, this study demonstrates that responsible collaboration between humans and AI serves as the key mechanism connecting Explainable AI with effective decision-making. The findings therefore extend the Explainable AI literature by highlighting that the effectiveness of AI-assisted decision-making depends not only on technical system characteristics but also on meaningful human involvement throughout the decision-making process.

Despite these contributions, this study has several limitations. The proposed model was evaluated using cross-sectional survey data from users of AI-enabled Intelligent Decision Support Systems, which may limit the generalizability of the findings across different organizational contexts and AI application domains. In addition, the relatively modest explanatory power of the structural model suggests that Decision Quality is also influenced by other factors that were not included in this study. Future research is therefore encouraged to examine additional antecedents such as AI trust, AI literacy, human oversight, perceived fairness, algorithmic accountability, and organizational AI governance. Longitudinal studies, experimental research designs, and cross-industry comparisons may also provide deeper insights into how responsible Human-AI Interaction evolves over time and contributes to sustainable AI-supported decision-making.

7. DECLARATIONS

7.1. About Authors

Cut Amalia Saffiera (CA)  <https://orcid.org/0009-0006-1693-5814>

Takumi Sase (TS)  <https://orcid.org/0000-0003-0240-946X>

Qurotul Aini (QA)  <https://orcid.org/0000-0002-7546-5721>

Danny Manongga (DM)  <https://orcid.org/0000-0002-7430-8740>

Harry Agustian (HA)  <https://orcid.org/0009-0006-8012-2260>

7.2. Author Contributions

Conceptualization: HA; Methodology: QA; Software: TS; Validation: DM and CA; Formal Analysis: QA and DM; Investigation: TS; Resources: CA; Data Curation: DM; Writing Original Draft Preparation: QA and TS; Writing Review and Editing: DM and CA; Visualization: HA; All authors, CA, TS, QA, DM, and HA, 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] A. Qazi and S. Ullah, "Human-ai collaboration in intelligent decision support systems," *Multidisciplinary Research in Computing Information Systems*, vol. 5, no. 9, pp. 691–702, 2025.
- [2] H. Nurhaeni, S. Kosasi, M. B. Thalia, E. A. Natalia, and J. Edwards, "Sentiment aware chatbots as companions for reducing loneliness through positive computing," *Journal of Orange Technology*, vol. 1, no. 1, pp. 39–50, 2024.
- [3] S. Baker and W. Xiang, "Explainable ai is responsible ai: How explainability creates trustworthy and socially responsible artificial intelligence," *arXiv preprint arXiv:2312.01555*, 2023.
- [4] 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.
- [5] K. Baum, S. Mantel, E. Schmidt, and T. Speith, "From responsibility to reason-giving explainable artificial intelligence," *Philosophy & Technology*, vol. 35, no. 1, p. 12, 2022.
- [6] V. Chamola, V. Hassija, A. R. Sulthana, D. Ghosh, D. Dhingra, and B. Sikdar, "A review of trustworthy and explainable artificial intelligence (xai)," *IEEE Access*, vol. 11, pp. 78 994–79 015, 2023.
- [7] H. Handoko, S. B. Waluya, W. Wardono, S. Sugiman, and L. W. Ming, "Geometry learning pathways for creative thinking via edupreneurial approaches," *Aptisi Transactions on Technopreneurship (ATT)*, vol. 8, no. 2, pp. 669–681, 2026.
- [8] U. Ehsan, P. Wintersberger, Q. V. Liao, E. A. Watkins, C. Manger, H. Daumé III, A. Riener, and 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.
- [9] J. Qadir, M. Q. Islam, and A. Al-Fuqaha, "Toward accountable human-centered ai: rationale and promising directions," *Journal of Information, Communication and Ethics in Society*, vol. 20, no. 2, pp. 329–342, 2022.
- [10] K. Kalasampath, K. Spoorthi, S. Sajeev, S. S. Kuppa, K. Ajay, and A. Maruthamuthu, "A literature review on applications of explainable artificial intelligence (xai)," *IEEE access*, vol. 13, pp. 41 111–41 140, 2025.
- [11] 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.
- [12] D. E. Mathew, D. U. Ebem, A. C. Ikegwu, P. E. Ukeoma, and N. F. Dibiaezue, "Recent emerging techniques in explainable artificial intelligence to enhance the interpretable and understanding of ai models for human: De mathew et al." *Neural processing letters*, vol. 57, no. 1, p. 16, 2025.
- [13] C. Meske, B. Abedin, M. Klier, and F. Rabhi, "Explainable and responsible artificial intelligence: Explainable and responsible artificial intelligence," *Electronic Markets*, vol. 32, no. 4, pp. 2103–2106, 2022.
- [14] 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.
- [15] P. N. Mahalle, R. V. Patil, N. Dey, R. G. Crespo, R. S. Sherratt *et al.*, "Explainable ai for human-centric ethical iot systems," *IEEE Transactions on Computational Social Systems*, vol. 11, no. 3, pp. 3407–3419, 2023.
- [16] M. Ren, N. Chen, and H. Qiu, "Human-machine collaborative decision-making: An evolutionary roadmap based on cognitive intelligence," *International Journal of Social Robotics*, vol. 15, no. 7, pp. 1101–1114, 2023.

- [17] 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.
- [18] M. A. Rahman and A. Rahman, "A systematic review of intelligent support systems for strategic decision-making using human-ai interaction in enterprise platforms," *Available at SSRN 5269458*, 2025.
- [19] M. Hardini, M. Siahaan, I. N. Hikam, T. Nurhaeni, and K. Vaher, "Examining the effects of innovation access safety and sustainability practices on global health outcomes," *Journal of Orange Technology*, vol. 1, no. 2, pp. 103–114, 2025.
- [20] 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.
- [21] G. Kostopoulos, G. Davrazos, and S. Kotsiantis, "Explainable artificial intelligence-based decision support systems: A recent review," *Electronics*, vol. 13, no. 14, p. 2842, 2024.
- [22] S. Gupta, S. Modgil, S. Bhattacharyya, and I. Bose, "Artificial intelligence for decision support systems in the field of operations research: review and future scope of research," *Annals of Operations Research*, vol. 308, no. 1, pp. 215–274, 2022.
- [23] R. T. H. Safariningsih, U. Rahardja, M. T. D. S. Putri, N. Azizah, E. Harris *et al.*, "Advanced big data analytics for proactive cyber threat mitigation in large scale computer networks," *Journal of Computer Science and Technology Application*, vol. 3, no. 2, pp. 199–212, 2026.
- [24] M. Elhaddad and S. Hamam, "Ai-driven clinical decision support systems: an ongoing pursuit of potential," *Cureus*, vol. 16, no. 4, p. e57728, 2024.
- [25] Q. Xu, W. Xie, B. Liao, C. Hu, L. Qin, Z. Yang, H. Xiong, Y. Lyu, Y. Zhou, and A. Luo, "Interpretability of clinical decision support systems based on artificial intelligence from technological and medical perspective: a systematic review," *Journal of healthcare engineering*, vol. 2023, no. 1, p. 9919269, 2023.
- [26] L. Meria, M. S. Gunawan, S. Solahudin, U. Rahardja, and A. Patel, "Enhancing digital business students competence through ui/ux design training for digital product development," *ADI Pengabdian Kepada Masyarakat*, vol. 6, no. 2, pp. 133–144, 2026.
- [27] S. Bayer, H. Gimpel, and M. Markgraf, "The role of domain expertise in trusting and following explainable ai decision support systems," *Journal of Decision Systems*, vol. 32, no. 1, pp. 110–138, 2022.
- [28] J. Amann, D. Vetter, S. N. Blomberg, H. C. Christensen, M. Coffee, S. Gerke, T. K. Gilbert, T. Hagendorff, S. Holm, M. Livne *et al.*, "To explain or not to explain?—artificial intelligence explainability in clinical decision support systems," *PLOS Digital Health*, vol. 1, no. 2, p. e0000016, 2022.
- [29] F. Shwede and H. M. Alzoubi, "Artificial intelligence (ai) integration into the decision support systems of health care centers," in *International Scientific Conference Management and Engineering*. Springer, 2024, pp. 165–176.
- [30] A. Hussain and K. Ullah, "An intelligent decision support system for spherical fuzzy sugeno-weber aggregation operators and real-life applications," *Spectrum of mechanical engineering and operational research*, vol. 1, no. 1, pp. 177–188, 2024.
- [31] S. Alon-Barkat and M. Busuioc, "Human–ai interactions in public sector decision making: "automation bias" and "selective adherence" to algorithmic advice," *Journal of Public Administration Research and Theory*, vol. 33, no. 1, pp. 153–169, 2023.
- [32] I. J. Akpan, Y. M. Kobara, J. Owolabi, A. A. Akpan, and O. F. Offodile, "Conversational and generative artificial intelligence and human–chatbot interaction in education and research," *International Transactions in Operational Research*, vol. 32, no. 3, pp. 1251–1281, 2025.
- [33] G. Lorenzini, L. Arbelaez Ossa, D. M. Shaw, and B. S. Elger, "Artificial intelligence and the doctor–patient relationship expanding the paradigm of shared decision making," *Bioethics*, vol. 37, no. 5, pp. 424–429, 2023.
- [34] "Explainable ai in clinical decision support systems: A meta-analysis of methods, applications, and usability challenges," *Healthcare*, 2025. [Online]. Available: <https://pmc.ncbi.nlm.nih.gov/articles/PMC12427955/>