

Developing a Qualitative Evaluation Framework for Explainable Decision Support Systems from a User Experience Perspective

Bernadus Gunawan Sudarsono^{1*}, Prima Dina Atika², Alexius Ulan Bani³
Bhayangkara Jakarta Raya University, Indonesia^{1,2}
Bung Karno University, Indonesia³

Corresponding Author: Bernadus Gunawan : Bernadus.gs@dsn.ubharajaya.ac.id

ARTICLE INFO

Keywords: Explainable Decision Support Systems, User Experience, Explainable Artificial Intelligence, Qualitative Evaluation Framework, Transparency.

Received : 25, May

Revised : 15, June

Accepted: 07, July

©2026 Sudarsono, Atika, Bani (s):

This is an open-access article distributed under the terms of the

[Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).



ABSTRACT

The growing adoption of Explainable Decision Support Systems (EDSS) across various decision-making domains has created a need for evaluation mechanisms that assess not only system performance but also the quality of user experience in understanding and trusting system recommendations. This study aims to develop a qualitative evaluation framework for EDSS from a user experience perspective by identifying the key factors influencing users' perceptions of transparency, understandability, trust, and system usefulness. A qualitative exploratory approach was employed in this research. Data were collected through in-depth interviews and focus group discussions involving 25 participants consisting of system users, developers, and user experience experts. The collected data were analyzed using thematic analysis to identify the principal themes relevant to EDSS evaluation. The findings reveal that explanation transparency, information relevance, ease of interpretation, and user control are the primary dimensions affecting user experience and system acceptance. Based on these findings, a structured qualitative evaluation framework was developed to assess the quality of explainability in decision support systems. The proposed framework contributes to the advancement of Explainable Artificial Intelligence research and provides practical guidance for developers in designing systems that are more transparent, understandable, and user-centered.

INTRODUCTION

Artificial Intelligence (AI) has become a fundamental technology in supporting decision-making processes across various sectors, including healthcare, finance, education, and public administration. The increasing use of AI-driven Decision Support Systems (DSS) has significantly improved efficiency and accuracy in processing large-scale data and generating recommendations. However, many AI models operate as "black boxes," making it difficult for users to understand how decisions are produced. This lack of transparency has raised concerns regarding trust, accountability, and acceptance of AI-based systems, particularly in high-stakes environments. Consequently, Explainable Artificial Intelligence (XAI) has emerged as an important research area aimed at improving the interpretability and transparency of AI-supported decision-making processes (Rong et al., 2024; Naveed et al., 2024).

Globally, the demand for explainable and trustworthy AI systems continues to increase as organizations seek to balance predictive performance with human-centered decision-making. Recent studies indicate that users are more likely to trust and adopt AI systems when explanations are understandable, relevant, and aligned with their cognitive needs (Kim et al., 2024). Furthermore, the integration of explainability into Decision Support Systems has been recognized as a critical factor in ensuring transparency, fairness, and user satisfaction (Khan et al., 2024). Despite rapid advancements in explainability techniques, practical implementation remains challenging because users often struggle to interpret complex explanations generated by AI models. Therefore, evaluating explainability from a user experience perspective has become an increasingly important issue in contemporary informatics research.

Several previous studies have investigated the evaluation of Explainable Artificial Intelligence from different perspectives. Hoffman et al. (2023) proposed various measurement dimensions for XAI, including trust, satisfaction, mental models, and human-AI performance. Similarly, Bernardo and Seva (2023) emphasized the importance of user-centered design in improving the quality of explainable AI interactions. More recently, Al-Ansari et al. (2024) reviewed user-centered evaluation approaches and found considerable variation in evaluation methods across studies. These findings suggest that although explainability has received significant scholarly attention, there remains no widely accepted framework capable of comprehensively evaluating explainable decision support systems from the perspective of user experience.

A more specific research gap can be identified in the lack of qualitative frameworks that integrate transparency, understandability, trust, usefulness, and user control into a unified evaluation model. Arrieta et al. (2020) reported that only a limited number of studies employed standardized evaluation frameworks, resulting in inconsistencies across empirical findings. Likewise, Naveed et al. (2024) highlighted the absence of comprehensive user-centered guidelines for evaluating explainable AI systems. Existing research predominantly focuses on technical performance and algorithmic interpretability rather than exploring how users experience and perceive explanations during decision-making activities. As a result, there is a need for a structured qualitative

evaluation framework capable of capturing the complexity of user interactions with explainable decision support systems.

Based on the identified research gap, this study aims to develop a qualitative evaluation framework for Explainable Decision Support Systems from a user experience perspective. The study focuses on identifying the key dimensions that influence users' perceptions of transparency, trust, understandability, usefulness, and control when interacting with explainable systems. Through a qualitative approach, the research seeks to explore how users interpret explanations and how these explanations affect system acceptance and decision-making confidence. Furthermore, the study intends to establish a set of evaluation criteria that can be systematically applied across different application domains. The resulting framework is expected to provide a more comprehensive understanding of human-centered explainability in decision support environments.

The contribution of this research is twofold. Theoretically, it enriches the literature on Explainable Artificial Intelligence, Human-Computer Interaction, and User Experience by proposing an integrated qualitative framework for evaluating explainability in decision support systems. Practically, the framework offers guidance for system developers, designers, and organizations in designing AI systems that are more transparent, understandable, and trustworthy for end users. The framework may also assist policymakers and practitioners in assessing the effectiveness of explainability features before deploying AI systems in real-world settings. Ultimately, the study supports the development of human-centered AI technologies that promote responsible and sustainable adoption of intelligent decision support systems.

THEORETICAL FRAMEWORK

Explainable Artificial Intelligence in Decision Support Systems

The rapid advancement of Artificial Intelligence (AI) has transformed decision-making processes across various domains, including healthcare, finance, transportation, and public administration. Decision Support Systems (DSS) increasingly rely on sophisticated machine learning algorithms to generate predictions and recommendations from large volumes of data. While these technologies improve decision accuracy and efficiency, many AI models remain opaque to end users, creating concerns regarding transparency and accountability. This phenomenon, commonly referred to as the "black-box problem," has become a significant barrier to user acceptance and trust in AI-driven systems. Explainable Artificial Intelligence (XAI) has emerged as a promising approach to address this challenge by providing understandable explanations of how AI systems generate their outputs (Mohseni et al., 2021).

Explainability is particularly important in decision support environments where users are expected to rely on system-generated recommendations. According to Adadi and Berrada (2020), explanations enable users to validate system outputs, understand underlying reasoning processes, and make more informed decisions. Furthermore, explainability contributes to ethical AI deployment by enhancing transparency, fairness, and accountability. As a result,

researchers increasingly recognize explainability as a critical requirement for effective human-AI collaboration. The success of a Decision Support System therefore depends not only on predictive performance but also on the ability of users to comprehend and evaluate system recommendations.

User Experience as a Foundation for Explainable Systems

User Experience (UX) has become an essential perspective in evaluating the effectiveness of Explainable AI systems. UX encompasses users' perceptions, emotions, and responses that arise during interactions with technology. Within the context of XAI, user experience extends beyond usability to include factors such as trust, satisfaction, transparency, and perceived usefulness. Research conducted by Ehsan et al. (2021) demonstrated that explanations significantly influence users' confidence in AI-assisted decision-making processes. Users are more likely to accept recommendations when explanations are understandable and aligned with their expectations.

Moreover, explanation quality is highly dependent on user characteristics, including expertise, prior knowledge, and contextual needs. Different users may require different forms of explanations to achieve the same level of understanding. For example, domain experts often prefer detailed technical explanations, whereas novice users may benefit from simplified and intuitive explanations. This variability highlights the importance of adopting a human-centered approach to explainability evaluation. Consequently, user experience has become a central dimension in assessing the effectiveness of Explainable Decision Support Systems (Chazette & Schneider, 2022).

Trust, Transparency, and Understandability in Explainable AI

Trust is widely recognized as one of the most important determinants of successful human-AI interaction. In AI-supported decision-making environments, users must be willing to rely on recommendations generated by intelligent systems. Trust is influenced by several factors, including system reliability, transparency, explanation quality, and perceived competence. According to Suresh, Lao, and Liccaldi (2021), transparent explanations help users develop appropriate levels of trust by enabling them to understand the rationale behind AI-generated decisions. When users perceive explanations as meaningful and relevant, they are more likely to consider the system trustworthy.

Transparency is closely associated with understandability, which refers to the extent to which users can comprehend system behavior and decision logic. Research by Bućinca, Malaya, and Gajos (2021) found that explanations can improve user understanding and calibration of trust when appropriately designed. However, excessive or overly technical explanations may overwhelm users and reduce their ability to interpret information effectively. Therefore, successful explainability requires a balance between completeness and comprehensibility. These findings suggest that trust, transparency, and understandability should be considered core dimensions within any evaluation framework for Explainable Decision Support Systems.

Human-Centered Evaluation of Explainable AI

Traditional evaluation methods for AI systems primarily focus on technical metrics such as accuracy, precision, recall, and model fidelity. Although

these indicators provide valuable information about system performance, they often fail to capture how users perceive and interact with explanations. Recent studies have emphasized the importance of human-centered evaluation approaches that assess explainability from the perspective of end users rather than solely from algorithmic performance measures (Liao & Varshney, 2021). Such approaches acknowledge that explainability is ultimately a human-centered concept that depends on users' cognitive and contextual needs.

Human-centered evaluation typically involves qualitative methods such as interviews, focus groups, think-aloud protocols, and observational studies. These methods allow researchers to explore users' perceptions, interpretations, and experiences in greater depth. Zhou, Holstein, and Vaughan (2023) argued that qualitative approaches are particularly effective for uncovering nuanced insights regarding user trust, understanding, and satisfaction. Despite growing interest in human-centered explainability, there remains a lack of standardized frameworks capable of systematically evaluating user experiences across different application domains. This limitation creates opportunities for developing more comprehensive qualitative evaluation models.

Qualitative Evaluation Frameworks for Explainable Decision Support Systems

A framework serves as a structured guide for identifying and assessing factors that influence system effectiveness. In the context of Explainable Decision Support Systems, evaluation frameworks are needed to ensure that explanations support users' cognitive processes and decision-making activities. Existing evaluation approaches often focus on isolated dimensions such as trust or interpretability, resulting in fragmented assessments of explainability. According to Chazette and Schneider (2022), the absence of integrated evaluation criteria limits the comparability of findings across studies and hinders the development of user-centered explainable systems.

A qualitative evaluation framework offers an opportunity to address these limitations by incorporating multiple dimensions of user experience into a single model. Such a framework can integrate transparency, understandability, trust, usefulness, user control, and satisfaction as interconnected evaluation components. By utilizing qualitative methods such as thematic analysis, researchers can capture rich and contextualized insights into how users experience explanations during decision-making processes. Therefore, the development of a qualitative evaluation framework contributes not only to the advancement of Explainable AI research but also to the practical design of more transparent, understandable, and user-centered decision support systems.

METHODS

Research Design and Approach

This study employed a qualitative research approach using an exploratory design to develop a qualitative evaluation framework for Explainable Decision Support Systems (EDSS) from a user experience perspective. A qualitative approach was considered appropriate because the study aimed to explore participants' perceptions, experiences, and interpretations regarding the explainability of decision support systems in real-world contexts. Exploratory

qualitative research enables researchers to investigate complex social and technological phenomena that are not yet fully understood and to generate conceptual frameworks grounded in empirical evidence (Creswell & Poth, 2024). Furthermore, the human-centered nature of explainable artificial intelligence requires in-depth investigation of user experiences, making qualitative inquiry particularly suitable for this research (Miller, 2023).

The study adopted a framework development design consisting of exploration, conceptualization, and validation stages. Initially, relevant dimensions of explainability and user experience were identified through an extensive review of contemporary literature. Subsequently, empirical data were collected from participants to explore factors influencing transparency, trust, understandability, usefulness, and user control within EDSS environments. The findings were then synthesized into a preliminary evaluation framework and refined through expert validation. This design is consistent with qualitative framework development studies that emphasize iterative theory building and stakeholder engagement (Tracy, 2020).

Population and Sampling Technique

The population of this study comprised individuals who had experience interacting with AI-based decision support systems in professional, academic, or organizational settings. Because the research focused on obtaining rich and relevant insights rather than statistical generalization, a non-probability sampling strategy was employed. Specifically, purposive sampling was used to select participants who possessed sufficient knowledge and practical experience related to explainable systems and user interaction with intelligent technologies. Purposive sampling is widely recommended in qualitative studies where participants are selected based on their ability to provide information relevant to the research objectives (Campbell et al., 2020).

A total of 25 participants were recruited, consisting of 12 system users, 8 software developers, and 5 user experience specialists. The sample size was determined based on the principle of information power, which suggests that a smaller sample may be sufficient when participants possess highly relevant expertise and when the study aims are narrowly focused (Malterud et al., 2021). This participant composition enabled the study to capture diverse perspectives regarding explainability, usability, trust, and transparency within decision support systems. The inclusion of multiple stakeholder groups also strengthened the credibility and comprehensiveness of the proposed evaluation framework.

Data Collection Techniques and Research Instruments

Data were collected through semi-structured interviews and focus group discussions. Semi-structured interviews allowed participants to express their experiences and perceptions freely while ensuring consistency across interview sessions. The interview protocol was developed based on dimensions commonly discussed in explainable artificial intelligence and user experience literature, including transparency, trust, understandability, usefulness, and user control (Meske et al., 2022). Focus group discussions were subsequently conducted to facilitate collective reflection and explore similarities and differences among participant perspectives.

The interview guide consisted of open-ended questions adapted from previous human-centered explainability studies and refined according to the objectives of the present research. To ensure content validity, the instrument was reviewed by three experts in artificial intelligence, human-computer interaction, and qualitative research methodology. Pilot interviews involving three participants were also conducted to evaluate clarity, relevance, and comprehensibility of the questions. Since qualitative research emphasizes credibility and trustworthiness rather than statistical reliability, methodological rigor was established through member checking, peer debriefing, and triangulation of data sources (Stahl & King, 2020).

Research Procedure

The research procedure consisted of five sequential stages. The first stage involved conducting a comprehensive literature review to identify theoretical foundations and existing evaluation dimensions related to Explainable Artificial Intelligence and user experience. The second stage involved designing the interview protocol and obtaining expert feedback regarding the appropriateness of the research instruments. Ethical approval and participant consent were obtained before commencing data collection. Participants were then recruited based on predefined inclusion criteria.

The third stage involved conducting interviews and focus group discussions, which were audio-recorded and transcribed verbatim. During the fourth stage, researchers performed preliminary coding and thematic categorization to identify emerging concepts related to explainability evaluation. The final stage involved synthesizing the findings into a conceptual framework and validating the framework through expert review. Feedback obtained during the validation process was incorporated to refine the framework and enhance its applicability in diverse decision support system contexts.

Data Analysis Techniques

The collected data were analyzed using thematic analysis following the six-phase procedure proposed by Braun and Clarke (2022), consisting of familiarization with data, generation of initial codes, searching for themes, reviewing themes, defining and naming themes, and producing the final report. Thematic analysis was selected because it provides a systematic and flexible approach for identifying patterns and meanings within qualitative datasets. This method is particularly effective for examining user experiences and perceptions in technology-related research.

To facilitate data organization and coding, NVivo 14 software was employed as a qualitative data analysis tool. NVivo enabled efficient management of interview transcripts, coding structures, thematic relationships, and analytical memos. The credibility of findings was enhanced through investigator triangulation, whereby multiple researchers independently reviewed coding results and discussed discrepancies until consensus was achieved. The resulting themes formed the basis for constructing the qualitative evaluation framework for Explainable Decision Support Systems from a user experience perspective.

RESULTS

Participant Characteristics

This study involved 25 participants selected through purposive sampling based on their experience with AI-based decision support systems. The participants consisted of 12 system users, 8 software developers, and 5 user experience specialists. The diversity of participants was intended to capture multiple perspectives regarding explainability, transparency, and user interaction with Explainable Decision Support Systems (EDSS). The collected data provided a comprehensive understanding of how different stakeholder groups perceive and evaluate explainability features within decision support environments.

Tabel 1.1 Participant Characteristics

Participant Group	Number of Participants	Percentage
System Users	12	48%
Software Developers	8	32%
User Experience Specialists	5	20%
Total	25	100%

Table 1.1 presents the demographic composition of the study participants. System users constituted the largest group, representing 48% of the total participants, followed by software developers (32%) and user experience specialists (20%). The predominance of system users enabled the study to capture direct experiences and perceptions regarding the usability and explainability of decision support systems. Meanwhile, the inclusion of developers and UX specialists provided complementary insights concerning system design, implementation challenges, and user-centered evaluation practices. The diversity of participants contributed to a more comprehensive understanding of explainability requirements from multiple stakeholder perspectives.

Identification of Key Dimensions of Explainability

Thematic analysis identified four dominant dimensions influencing user experience in Explainable Decision Support Systems: transparency, information relevance, ease of interpretation, and user control. These dimensions consistently emerged across all participant groups and were frequently associated with positive perceptions of system effectiveness.

Participants emphasized that transparency plays a critical role in helping users understand the reasoning behind system recommendations. Most participants indicated that explanations should clearly describe the factors influencing decisions rather than merely presenting final outcomes. In addition, users expressed a preference for explanations that provide contextual information relevant to their specific tasks and decision-making objectives.

Tabel 1.2 Main Dimensions Identified Through Thematic Analysis

Dimension	Description
Transparency	Clarity regarding how recommendations are generated
Information Relevance	Alignment of explanations with user needs and decision contexts
Ease of Interpretation	Simplicity and understandability of explanations

Dimension	Description
User Control	Ability to explore, modify, or request additional explanations

Table 1.2 summarizes the four primary dimensions identified through thematic analysis. These dimensions emerged consistently across interviews and focus group discussions and were frequently associated with positive user experiences. Transparency was considered essential for understanding how recommendations were generated, while information relevance ensured that explanations were meaningful within specific decision-making contexts. Ease of interpretation emphasized the importance of presenting explanations in an understandable format, whereas user control reflected the need for interactive and flexible explanatory features. Collectively, these dimensions formed the conceptual foundation of the proposed qualitative evaluation framework.

User Perceptions of Transparency

Analysis revealed that participants considered transparency as the foundational component of explainability. Users reported greater confidence when systems explicitly communicated the reasoning process used to generate recommendations. Conversely, systems providing only recommendations without explanatory support were frequently perceived as difficult to trust.

Several participants noted that transparency reduced uncertainty during decision-making activities. Explanations that identified influential variables, data sources, and reasoning pathways enabled users to better assess the validity of recommendations. Participants also indicated that transparent explanations facilitated learning and increased familiarity with system functionality over time.

Table 1.3 User Responses Regarding Transparency

Observation	Frequency of Mention
Increased trust in recommendations	High
Improved understanding of decisions	High
Reduced uncertainty	Moderate
Enhanced learning experience	Moderate

As shown in Table 1.3, transparency was strongly associated with increased trust and improved understanding of system recommendations. Participants frequently reported that transparent explanations enabled them to comprehend the reasoning process underlying AI-generated outputs, thereby enhancing confidence in the system. Transparency also contributed to reducing uncertainty during decision-making activities by allowing users to verify the validity of recommendations. Furthermore, several participants indicated that transparent explanations facilitated learning and helped them better understand system behavior over time. These findings suggest that transparency is a critical determinant of explainability effectiveness.

Importance of Information Relevance and Interpretability

The analysis further revealed that the usefulness of explanations depends not only on transparency but also on their relevance to the user's immediate

decision context. Participants consistently indicated that explanations should directly address the rationale underlying specific recommendations rather than providing generic descriptions of system operations.

Ease of interpretation emerged as another critical factor. Users preferred concise explanations written in clear language and organized in a logical structure. Participants frequently reported difficulties understanding highly technical explanations, particularly when domain-specific terminology was used without sufficient clarification.

Table 1.4 Factors Influencing Explanation Effectiveness

Factor	Impact on User Experience
Relevant information	Very High
Clear language	Very High
Contextual explanation	High
Technical detail	Moderate
Excessive complexity	Negative

Table 1.4 illustrates the factors influencing the effectiveness of explanations from the user perspective. Relevant information and clear language were identified as the most influential factors because they enabled users to quickly understand and apply system recommendations. Contextual explanations were also considered important as they connected system outputs to specific decision-making scenarios. In contrast, highly technical explanations produced mixed responses; while some expert users appreciated detailed information, many participants found excessive technical complexity difficult to understand. Consequently, explanation design should balance completeness with accessibility to accommodate users with varying levels of expertise.

User Control and Interactive Explainability

Another significant finding concerned the role of user control in explainable systems. Participants expressed a preference for systems that allow users to explore explanations at varying levels of detail according to their individual needs. Rather than receiving static explanations, users favored interactive mechanisms that support further inquiry and clarification.

Developers and UX specialists particularly emphasized the importance of adaptive explainability features. They suggested that systems should provide flexible explanation interfaces capable of accommodating users with different levels of expertise. Novice users often required simplified explanations, whereas experienced users preferred access to more detailed information.

Table 1.5 Preferred Explainability Features

Feature	Frequency
Expandable explanations	High
Interactive exploration	High
Multiple explanation levels	High
Visual explanation support	Moderate
Static explanations only	Low

Table 1.5 presents participants' preferences regarding explainability features. The results indicate a strong preference for interactive explanation mechanisms, particularly expandable explanations, interactive exploration tools, and multiple levels of explanation detail. Participants valued the ability to access additional information when necessary while maintaining concise explanations during routine use. Visual explanation support was also viewed positively, although it was mentioned less frequently than interactive features. Conversely, static explanations received limited support because participants perceived them as inflexible and insufficient for addressing diverse information needs. These findings highlight the importance of adaptive and user-centered explanation design.

Development of the Qualitative Evaluation Framework

Based on the identified themes, a qualitative evaluation framework for Explainable Decision Support Systems was developed. The framework integrates the four primary dimensions identified during data analysis and organizes them into a structured assessment model. Each dimension contains several evaluation indicators derived directly from participant experiences and observations.

Table 1.6 Proposed Framework Dimensions

Dimension	Evaluation Focus
Transparency	Visibility of reasoning processes
Information Relevance	Alignment with decision context
Ease of Interpretation	Clarity and comprehensibility
User Control	Flexibility and interaction with explanations

Table 1.6 presents the final dimensions incorporated into the proposed qualitative evaluation framework. Transparency focuses on the extent to which users can understand the reasoning process underlying system recommendations. Information relevance assesses whether explanations provide meaningful information that supports decision-making activities. Ease of interpretation evaluates the clarity and comprehensibility of explanations, while user control measures the degree of flexibility available to users when interacting with explanatory features. Together, these dimensions provide a structured approach for evaluating Explainable Decision Support Systems from a user experience perspective and serve as the foundation for future framework implementation and validation studies.

DISCUSSION

Transparency as the Foundation of User Trust in Explainable Decision Support Systems

The findings indicate that transparency emerged as the most influential dimension affecting users' acceptance of Explainable Decision Support Systems (EDSS). Participants consistently reported that explanations describing how recommendations were generated increased their confidence in system outputs and reduced uncertainty during decision-making processes. These findings

support the theoretical perspective that transparency serves as a prerequisite for establishing trust in intelligent systems. According to Mohseni, Zarei, and Ragan (2021), transparency enables users to construct mental models of AI behavior, thereby facilitating more informed judgments regarding system reliability. Similarly, Ehsan, Liao, Muller, Riedl, and Weisz (2021) argued that explanations are most effective when they provide meaningful insight into the reasoning process rather than merely presenting outcomes.

From a theoretical standpoint, these results align with the principles of Human-Centered Explainable Artificial Intelligence, which emphasize that explainability should be evaluated based on users' understanding rather than algorithmic interpretability alone. The prominence of transparency observed in this study suggests that users value explanations that enable them to assess the validity and credibility of recommendations independently. This finding extends previous research by demonstrating that transparency functions not only as an informational attribute but also as a psychological mechanism that strengthens confidence and promotes acceptance of AI-supported decisions. Consequently, system developers should prioritize explanatory features that clearly communicate reasoning pathways, influential variables, and decision criteria.

The Importance of Information Relevance and Cognitive Alignment

Another significant finding concerns the role of information relevance in shaping positive user experiences. Participants emphasized that explanations must directly address the context of the decision being made rather than providing generic descriptions of system operations. This result supports Cognitive Fit Theory, which proposes that information becomes more useful when its presentation aligns with users' task requirements and cognitive processes (Vessey, 2021). Explanations that were perceived as relevant enabled participants to understand recommendations more efficiently and apply them more effectively in decision-making situations.

The findings are also consistent with the work of Chazette and Schneider (2022), who argued that explainability should be considered a user-centered requirement rather than merely a technical feature. Participants demonstrated a preference for contextual explanations that explicitly linked recommendations to specific circumstances and objectives. This suggests that explanation quality depends not only on the amount of information provided but also on its relevance to users' immediate needs. Theoretically, these results contribute to the growing body of literature emphasizing contextual explainability as a critical determinant of effective human-AI interaction. Practically, they indicate that future EDSS designs should incorporate adaptive explanation mechanisms capable of tailoring information to different decision contexts.

Ease of Interpretation and User-Centered Explainability

The study revealed that ease of interpretation significantly influences users' ability to understand and utilize explanations. Participants consistently preferred concise explanations expressed in clear language over highly technical descriptions. This finding aligns with Cognitive Load Theory, which suggests that excessive complexity can overload users' cognitive resources and hinder comprehension (Sweller, Van Merriënboer, & Paas, 2021). Participants reported

that overly detailed technical explanations often reduced the usefulness of the system despite providing more information.

These findings support recent studies emphasizing the importance of simplicity and usability in explainable systems. For example, Bućinca, Malaya, and Gajos (2021) found that explanations are most effective when they facilitate critical thinking without overwhelming users with unnecessary details. The results indicate that explainability should be viewed as a communication process rather than solely a technical capability. Therefore, effective EDSS designs should balance completeness with comprehensibility by presenting explanations in a format that accommodates users with varying levels of expertise. This finding contributes to explainable AI research by highlighting the importance of integrating user experience principles into explanation design.

User Control as a Dimension of Interactive Explainability

A notable contribution of this study is the identification of user control as a key dimension within the proposed evaluation framework. Participants expressed a strong preference for interactive explanations that allow users to access varying levels of detail according to their individual needs. This finding supports the concept of interactive explainability proposed by Abdul, Vermeulen, Wang, Lim, and Kankanhalli (2020), who argued that users should be able to engage dynamically with explanations rather than passively receiving static information.

The emphasis on user control suggests that explainability is not a one-directional process but rather an ongoing interaction between users and intelligent systems. Participants valued features such as expandable explanations, additional information requests, and customizable explanation depth. These preferences indicate that users seek autonomy in determining how much information they require to make informed decisions. Theoretically, this finding extends previous explainability frameworks by incorporating user agency as an essential component of evaluation. Practically, it highlights the importance of designing EDSS interfaces that support exploratory and adaptive explanation experiences.

Contribution of the Proposed Qualitative Evaluation Framework

The integration of transparency, information relevance, ease of interpretation, and user control into a single framework represents a significant contribution to explainable AI literature. Previous studies often examined these dimensions independently, resulting in fragmented evaluation approaches (Zhou, Holstein, and Vaughan, 2023). In contrast, the framework developed in this study provides a holistic perspective that captures multiple aspects of user experience simultaneously.

The framework contributes theoretically by extending Human-Computer Interaction and Explainable Artificial Intelligence research through the introduction of a user-centered qualitative evaluation model. Unlike conventional evaluation approaches that emphasize algorithmic performance metrics, the proposed framework focuses on how users experience and interpret explanations. Practically, the framework offers guidance for organizations and developers seeking to assess and improve explainability features in decision

support systems. The framework may also serve as a foundation for future studies investigating user-centered explainability across diverse application domains.

CONCLUSIONS AND RECOMMENDATIONS

This study developed a qualitative evaluation framework for Explainable Decision Support Systems (EDSS) from a user experience perspective. The findings revealed four key dimensions that significantly influence users' perceptions and acceptance of explainable systems, namely transparency, information relevance, ease of interpretation, and user control. Among these dimensions, transparency emerged as the most influential factor in fostering trust and confidence in AI-generated recommendations. The proposed framework contributes to the advancement of Explainable Artificial Intelligence and Human-Computer Interaction research by providing a structured, user-centered approach for evaluating explainability beyond traditional technical performance measures. Furthermore, the framework offers practical guidance for developers and organizations seeking to design more transparent, understandable, and user-oriented decision support systems.

Based on the findings, it is recommended that developers prioritize explanation mechanisms that are transparent, contextually relevant, easy to understand, and adaptable to different user needs. Organizations implementing AI-based decision support systems should also consider user experience evaluation as an integral component of system assessment and deployment. In addition, stakeholders should encourage the adoption of human-centered design principles to improve trust, usability, and acceptance of explainable technologies across various application domains.

FURTHER STUDY

Future research is encouraged to validate the proposed framework using quantitative and mixed-method approaches involving larger and more diverse participant populations. Comparative studies across different sectors, such as healthcare, finance, education, and public services, may provide further insights into the applicability and generalizability of the framework. Future investigations should also explore additional dimensions of explainability, including fairness, accountability, ethical transparency, and user empowerment, which are increasingly important in contemporary AI governance. Moreover, empirical implementation of the framework within operational decision support systems would provide valuable evidence regarding its effectiveness in real-world environments.

ACKNOWLEDGMENT

The authors would like to express their sincere gratitude to all participants who voluntarily contributed their time, experiences, and insights to this study. Appreciation is also extended to the experts and practitioners who provided valuable feedback during the framework development and validation processes. Their contributions were instrumental in enhancing the quality and relevance of

this research. Finally, the authors acknowledge the support of their affiliated institutions for facilitating the successful completion of this study.

REFERENCES

- Abdul, A., Vermeulen, J., Wang, D., Lim, B. Y., & Kankanhalli, M. (2018). Trends and trajectories for explainable, accountable and intelligible systems: An HCI research agenda. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems* (pp. 1–18). Association for Computing Machinery. <https://doi.org/10.1145/3173574.3174156>
- Adadi, A., & Berrada, M. (2018). Peeking inside the black-box: A survey on explainable artificial intelligence (XAI). *IEEE Access*, 6, 52138–52160. <https://doi.org/10.1109/ACCESS.2018.2870052>
- Al-Ansari, N., Al-Thani, D., & Al-Mansoori, R. S. (2024). User-centered evaluation of explainable artificial intelligence (XAI): A systematic literature review. *Human Behavior and Emerging Technologies*, 2024, Article 4628855. <https://doi.org/10.1155/2024/4628855>
- Barredo Arrieta, A., Díaz-Rodríguez, N., Del Ser, J., Bannetot, A., Tabik, S., Barbado, A., García, S., Gil-López, S., Molina, D., Benjamins, R., Chatila, R., & Herrera, F. (2020). Explainable artificial intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, 82–115. <https://doi.org/10.1016/j.inffus.2019.12.012>
- Bernardo, E., & Seva, R. (2023). Affective design analysis of explainable artificial intelligence (XAI): A user-centric perspective. *Informatics*, 10(1), Article 32. <https://doi.org/10.3390/informatics10010032>
- Braun, V., & Clarke, V. (2022). *Thematic analysis: A practical guide*. SAGE Publications.
- Buçinca, Z., Malaya, M. B., & Gajos, K. Z. (2021). To trust or to think: Cognitive forcing functions can reduce overreliance on AI in AI-assisted decision-making. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW1), Article 188, 1–21. <https://doi.org/10.1145/3449287>
- Campbell, S., Greenwood, M., Prior, S., Shearer, T., Walkem, K., Young, S., Bywaters, D., & Walker, K. (2020). Purposive sampling: Complex or simple? Research case examples. *Journal of Research in Nursing*, 25(8), 652–661. <https://doi.org/10.1177/1744987120927206>
- Chazette, L., Klös, V., Herzog, F., & Schneider, K. (2022). Requirements on explanations: A quality framework for explainability. In *2022 IEEE 30th International Requirements Engineering Conference (RE)* (pp. 140–152). IEEE. <https://doi.org/10.1109/RE54965.2022.00019>
- Creswell, J. W., & Poth, C. N. (2024). *Qualitative inquiry and research design: Choosing among five approaches* (5th ed.). SAGE Publications.
- Ehsan, U., Liao, Q. V., Muller, M., Riedl, M. O., & Weisz, J. D. (2021). Expanding explainability: Towards social transparency in AI systems. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems* (pp. 1–19). Association for Computing Machinery. <https://doi.org/10.1145/3411764.3445188>

- Liao, Q. V., & Varshney, K. R. (2021). Human-centered explainable AI (XAI): From algorithms to user experiences. *arXiv*. <https://doi.org/10.48550/arXiv.2110.10790>
- Malterud, K., Siersma, V. D., & Guassora, A. D. (2016). Sample size in qualitative interview studies: Guided by information power. *Qualitative Health Research*, 26(13), 1753–1760. <https://doi.org/10.1177/1049732315617444>
- Meske, C., Bunde, E., Schneider, J., & Gersch, M. (2022). Explainable artificial intelligence: Objectives, stakeholders, and future research opportunities. *Information Systems Management*, 39(1), 53–63. <https://doi.org/10.1080/10580530.2020.1849465>
- Miller, T. (2023). Explainable AI is dead, long live explainable AI! Hypothesis-driven decision support using evaluative AI. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency* (pp. 333–342). Association for Computing Machinery. <https://doi.org/10.1145/3593013.3594001>
- Mohseni, S., Zarei, N., & Ragan, E. D. (2021). A multidisciplinary survey and framework for design and evaluation of explainable AI systems. *ACM Transactions on Interactive Intelligent Systems*, 11(3–4), Article 24, 1–45. <https://doi.org/10.1145/3387166>
- Naveed, S., Stevens, G., & Robin-Kern, D. (2024). An overview of the empirical evaluation of explainable AI (XAI): A comprehensive guideline for user-centered evaluation in XAI. *Applied Sciences*, 14(23), Article 11288. <https://doi.org/10.3390/app142311288>
- Rong, Y., Leemann, T., Nguyen, T.-T., Fiedler, L., Qian, P., Unhelkar, V., Seidel, T., Kasneci, G., & Kasneci, E. (2024). Towards human-centered explainable AI: A survey of user studies for model explanations. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 46(4), 2104–2122. <https://doi.org/10.1109/TPAMI.2023.3331846>
- Stahl, N. A., & King, J. R. (2020). Expanding approaches for research: Understanding and using trustworthiness in qualitative research. *Journal of Developmental Education*, 44(1), 26–28.
- Suresh, H., Lao, N., & Liccardi, I. (2020). Misplaced trust: Measuring the interference of machine learning in human decision-making. In *Proceedings of the 12th ACM Conference on Web Science* (pp. 315–324). Association for Computing Machinery. <https://doi.org/10.1145/3394231.3397922>
- Sweller, J., van Merriënboer, J. J. G., & Paas, F. (2019). Cognitive architecture and instructional design: 20 years later. *Educational Psychology Review*, 31(2), 261–292. <https://doi.org/10.1007/s10648-019-09465-5>
- Tracy, S. J. (2020). *Qualitative research methods: Collecting evidence, crafting analysis, communicating impact* (2nd ed.). Wiley-Blackwell.
- Vessey, I. (1991). Cognitive fit: A theory-based analysis of the graphs versus tables literature. *Decision Sciences*, 22(2), 219–240. <https://doi.org/10.1111/j.1540-5915.1991.tb00344.x>