



Explainable Artificial Intelligence for Engineering Decision Support Systems

Rezi Munizar^{1*}, Ahmad Budi Trisnawan², Maiza Fikri³, Tarma⁴, Eko Sutrisno⁵

¹ Universitas Ratu Samban, Indonesia

² Universitas Mahakarya Asia, Indonesia

³ Institut Teknologi dan Bisnis Bina Sriwijaya Palembang, Indonesia

⁴ Universitas Negeri Jakarta, Indonesia

⁵ Universitas Islam Majapahit, Indonesia

email: rezimunizar@gmail.com¹

Article Info :

Received:

19-05-2026

Revised:

01-06-2026

Accepted:

06-06-2026

Abstract

Explainable Artificial Intelligence (XAI) has become increasingly important in Engineering Decision Support Systems because growing algorithmic complexity often reduces transparency, accountability, and practitioner confidence despite substantial improvements in predictive capability. This study develops a conceptual framework that integrates engineering data processing, artificial intelligence inference, explainability mechanisms, and human-centered decision support within a unified multilayer architecture. A non-empirical system design methodology based on design-oriented systems research, architecture-based evaluation, and scenario-driven analytical simulation is employed to examine architectural consistency, interoperability, traceability, and explainability across representative engineering contexts. The analytical results indicate that explanation fidelity, interpretability, transparency, traceability, and modular scalability function as complementary engineering quality attributes that collectively strengthen trustworthy decision support while preserving logical consistency between engineering evidence and computational reasoning. The proposed architecture also demonstrates technology independence and adaptability across heterogeneous engineering domains through explicit modular interactions and standardized information flows. This study contributes an architecture-centered perspective that advances theoretical understanding of explainable engineering intelligence while providing a reproducible conceptual foundation for future empirical implementation, quantitative validation, and standardized evaluation of trustworthy Engineering Decision Support Systems.

Keywords : Explainable Artificial Intelligence, Engineering Decision Support Systems, Human-Centered Intelligence, Transparent Decision Making, System Architecture



©2022 Authors.. This work is licensed under a Creative Commons Attribution-Non Commercial 4.0 International License.
(<https://creativecommons.org/licenses/by-nc/4.0/>)

INTRODUCTION

The rapid evolution of artificial intelligence (AI) has fundamentally transformed engineering decision support systems (DSSs), shifting them from conventional rule-based frameworks toward adaptive, data-driven architectures capable of processing complex, multidimensional, and dynamic operational environments. Recent advances in deep learning, reinforcement learning, graph neural networks, and foundation models have substantially improved predictive accuracy across engineering domains, including manufacturing, construction, energy management, cyber-physical systems, infrastructure maintenance, and industrial automation. Despite these remarkable advances, increasing model complexity has simultaneously intensified concerns regarding algorithmic opacity, accountability, and the inability of decision-makers to interpret machine-generated recommendations within safety-critical engineering contexts where erroneous decisions may result in significant economic losses, operational failures, or public safety risks. Explainable Artificial Intelligence (XAI) has consequently emerged as a central research paradigm seeking to reconcile predictive performance with interpretability by providing transparent reasoning mechanisms capable of enhancing users' cognitive understanding, trust, and confidence in AI-assisted decision processes. Contemporary engineering research increasingly recognizes explainability not merely as an auxiliary visualization feature but as an essential design principle that determines whether intelligent decision support systems

can satisfy technical, organizational, regulatory, and ethical expectations in real-world deployment environments (Hussain et al., 2021; Panagoulas et al., 2023; Kostopoulos et al., 2024; Ravi et al., 2025).

Existing scholarship has progressively demonstrated that explainability contributes to engineering decision quality through multiple complementary mechanisms, although these contributions remain distributed across domain-specific applications rather than consolidated into an integrated engineering decision framework. Studies within manufacturing system design have illustrated how structured information models facilitate traceable reasoning pathways between engineering knowledge and AI-generated recommendations, allowing designers to understand the rationale underlying optimization outcomes while preserving analytical consistency (Cochran et al., 2022). Research in energy management has extended this perspective by proposing tailored explainability methodologies that adapt explanation formats according to stakeholder characteristics, revealing that explanation effectiveness depends not only on algorithmic transparency but also on contextual communication strategies (Panagoulas et al., 2023). Investigations in construction engineering have introduced counterfactual explanations capable of supporting risk-based decision-making through alternative scenario generation that enhances practitioners' understanding of causal relationships among engineering variables (Zhan et al., 2024). Parallel developments in healthcare decision support have confirmed that integrating explainability into predictive models improves practitioners' confidence when evaluating machine-learning recommendations in high-stakes diagnostic environments (Khanna et al., 2023). Recent studies addressing cyber resilience have further emphasized that explainability strengthens agile decision-making under uncertainty by improving organizational responsiveness to rapidly evolving operational threats (Sadeghi et al., 2024). Collectively, these investigations suggest that XAI possesses substantial potential for engineering decision support, although the accumulated evidence reflects methodological diversity and conceptual fragmentation rather than theoretical convergence.

Despite this expanding body of knowledge, the current literature exhibits several unresolved limitations that restrict the maturation of Explainable Artificial Intelligence as a coherent engineering decision support paradigm. Existing investigations frequently evaluate explainability through isolated technical metrics without sufficiently examining whether different explanation mechanisms consistently improve engineering reasoning across heterogeneous operational environments. Domain-specific implementations remain largely disconnected, creating limited opportunities for transferring explainability principles between manufacturing, construction, energy systems, healthcare engineering, supply chains, and cyber-physical infrastructures despite sharing common decision-support characteristics (Cochran et al., 2022; Khanna et al., 2023; Sadeghi et al., 2024). Numerous studies continue emphasizing local interpretability techniques while paying comparatively little attention to the interaction among human cognition, organizational workflows, engineering expertise, and adaptive AI behavior during complex decision-making processes (Hussain et al., 2021; Kostopoulos et al., 2024). Emerging reviews acknowledge the rapid proliferation of explainability techniques but simultaneously highlight persistent inconsistencies concerning evaluation criteria, validation methodologies, explanation quality metrics, and stakeholder-centered performance assessment, making cross-study comparison increasingly difficult (Kostopoulos et al., 2024; Ravi et al., 2025). Recent engineering applications addressing threat modelling have likewise demonstrated promising operational capabilities while revealing insufficient theoretical integration between explainability mechanisms and engineering asset governance, thereby exposing an important conceptual discontinuity within current research trajectories (Takon, 2025).

These unresolved inconsistencies possess profound scientific and practical implications because engineering decisions increasingly occur within environments characterized by uncertainty, interdependent infrastructures, autonomous systems, and stringent accountability requirements that cannot tolerate opaque computational reasoning. Industrial organizations, infrastructure operators, policymakers, and engineering practitioners require decision support systems capable of delivering not only accurate predictions but also explanations that satisfy transparency, auditability, regulatory compliance, and collaborative human–AI interaction. Insufficient explainability reduces user trust, weakens organizational adoption, complicates responsibility attribution, and limits the operational value of advanced AI despite significant improvements in predictive performance. Contemporary engineering systems simultaneously demand explanation mechanisms that accommodate varying stakeholder expertise, diverse engineering objectives, and evolving operational contexts without

sacrificing computational efficiency or predictive reliability. Recent literature consistently argues that explainability should become an integral architectural component of engineering decision support rather than a post hoc interpretability layer, yet systematic understanding regarding how explainability can be embedded throughout engineering decision processes remains underdeveloped (Panagoulas et al., 2023; Zhan et al., 2024; Sadeghi et al., 2024; Takon, 2025; Ravi et al., 2025).

The present study positions itself within this unresolved scholarly landscape by addressing the conceptual separation between explainability techniques and engineering decision-support architecture through a comprehensive analytical perspective that synthesizes technological, methodological, and human-centered dimensions into a unified engineering framework. Rather than treating explainability as an isolated algorithmic property, this research conceptualizes Explainable Artificial Intelligence as a multidimensional capability that shapes decision quality through the interaction among model transparency, contextual reasoning, stakeholder interpretability, and engineering governance. Such positioning extends existing scholarship beyond application-specific demonstrations toward a broader theoretical understanding of how explainability contributes to reliable engineering decision support across heterogeneous domains characterized by varying operational complexities, risk profiles, and organizational constraints. This perspective seeks to establish stronger conceptual integration between AI explainability research and engineering decision science while clarifying the mechanisms through which transparent intelligence influences decision effectiveness, accountability, and sustainable technology adoption.

The objective of this research is to develop a comprehensive conceptual framework for Explainable Artificial Intelligence in engineering decision support systems by integrating technological explainability mechanisms, human-centered decision processes, and engineering governance principles into a coherent analytical model capable of supporting trustworthy and accountable AI-assisted decision-making. The study contributes theoretically by refining the conceptual boundaries of explainability within engineering decision science and proposing an integrative perspective that reconciles predictive intelligence with transparent reasoning. Methodologically, the research contributes a structured analytical framework for evaluating explainability across engineering decision contexts, providing a foundation for future empirical validation and interdisciplinary development of trustworthy intelligent engineering systems.

RESEARCH METHODS

This study adopts a non-empirical system design approach aimed at developing a conceptual framework for Explainable Artificial Intelligence (XAI) in Engineering Decision Support Systems (EDSS). The methodological foundation is based on design-oriented systems research, in which engineering decision-making is modeled as a multilayer architecture integrating data acquisition, predictive intelligence, explainability mechanisms, and human-centered decision interfaces. The proposed framework consists of four interconnected modules: (i) an engineering data processing layer responsible for feature extraction, normalization, and contextual representation; (ii) an AI inference layer that produces predictive or prescriptive decisions using machine learning models; (iii) an explainability layer incorporating both intrinsic and post-hoc explanation mechanisms, including feature attribution, rule extraction, counterfactual reasoning, and local-global interpretability strategies; and (iv) a decision support layer that translates computational explanations into engineering-relevant knowledge for different stakeholder groups. System modeling follows a modular architecture to ensure interoperability, extensibility, and technology independence, allowing the framework to be applied across manufacturing, construction, energy management, cyber-physical systems, and other engineering domains. Algorithmic interactions among these modules are formally specified through sequential information flow, ensuring that every decision generated by the predictive engine is accompanied by structured explanatory outputs capable of supporting transparent engineering reasoning. The complete architectural workflow is documented using standardized functional decomposition and process modeling to ensure methodological reproducibility and facilitate future implementation in practical engineering environments.

The technical analysis employs architecture-based evaluation and scenario-driven simulation to examine the consistency, scalability, and explainability of the proposed framework under representative engineering decision scenarios. Rather than validating predictive accuracy using domain-specific datasets, the framework is assessed through systematic analytical verification based on engineering

design principles and established XAI quality dimensions. Evaluation focuses on explanation fidelity, interpretability, transparency, consistency, traceability, modular scalability, computational feasibility, and decision support capability. Multiple engineering scenarios representing varying levels of decision complexity, uncertainty, and operational risk are analyzed to verify whether the proposed architecture maintains coherent information flow between predictive inference and explanation generation while preserving human interpretability. Internal validation is conducted through logical consistency analysis, architectural completeness assessment, and cross-layer dependency verification to ensure that every functional component satisfies its intended role without introducing structural redundancy or information ambiguity. Reproducibility is ensured through explicit specification of architectural components, module interactions, information pathways, evaluation criteria, and analytical procedures, enabling future researchers to replicate, implement, extend, or empirically validate the proposed Explainable Artificial Intelligence framework within diverse Engineering Decision Support System applications.

RESULTS AND DISCUSSION

Conceptual Architecture of Explainable Artificial Intelligence for Engineering Decision Support Systems

Engineering decision support systems increasingly require architectural designs that integrate predictive intelligence with transparent reasoning instead of treating explainability as an auxiliary visualization component. The conceptual framework developed in this study models Explainable Artificial Intelligence (XAI) as an interconnected architecture consisting of engineering data processing, AI inference, explainability generation, and human-centered decision support modules. Each module contributes a distinct functional responsibility while maintaining sequential information dependency that preserves semantic consistency throughout the decision-making process. Recent investigations have similarly argued that explainability should be embedded within system architecture because interpretability emerges from coordinated interactions among computational components rather than isolated explanation algorithms (Kostopoulos et al., 2024).

The analytical modeling indicates that the engineering data processing layer performs a foundational role by transforming heterogeneous engineering information into structured representations suitable for intelligent inference. Feature normalization, contextual encoding, uncertainty identification, and engineering knowledge representation collectively establish the informational quality required for reliable explanation generation. Predictive models operating on poorly structured engineering data frequently produce explanations that are internally coherent but operationally misleading because contextual dependencies remain hidden. Comparable observations have been reported in industrial engineering research emphasizing that information modeling determines the quality of explainable engineering recommendations before algorithmic inference is executed (Cochran et al., 2022; Ahmed et al., 2022).

The inference layer represents the computational core of the proposed architecture because predictive reasoning and explanatory reasoning are designed to evolve simultaneously rather than sequentially. Conventional decision support systems commonly prioritize predictive optimization while postponing interpretability until after model execution through external explanation techniques. The present conceptual model restructures this relationship by allowing explanation-aware inference pathways that preserve causal consistency between prediction and interpretation. Engineering perspectives on XAI have consistently recognized that integrated inference and explanation improve the reliability of intelligent systems operating within complex industrial environments (Hussain et al., 2021; Petrauskas et al., 2023).

The explainability layer extends beyond conventional feature importance visualization by incorporating complementary reasoning mechanisms that support engineering interpretation from multiple analytical perspectives. Counterfactual reasoning, local attribution, global model transparency, and rule-based explanation collectively produce richer engineering knowledge than single-method interpretability approaches. Different engineering scenarios demand different explanatory perspectives because operational objectives vary according to system complexity, stakeholder expertise, and acceptable risk thresholds.

Similar multidimensional explainability strategies have demonstrated improved decision transparency in engineering environments characterized by uncertainty and heterogeneous operational

constraints (Praveenraj et al., 2023; Zhan et al., 2024). The functional relationships among the four architectural modules become clearer when their engineering responsibilities are examined collectively. Table 1 summarizes the analytical decomposition developed through the conceptual architecture proposed in this study.

Table 1. Functional Architecture of the Proposed Explainable AI Framework

Architectural Module	Primary Function			Explainability Contribution	Engineering Outcome
Engineering Processing	Data	Feature extraction and contextual representation		Preserves semantic consistency	Reliable engineering information
AI Inference		Predictive and prescriptive analytics		Explanation-aware inference	Transparent prediction generation
Explainability Layer		Attribution, rule extraction, counterfactual reasoning		Human-understandable reasoning	Interpretable engineering knowledge
Decision Layer	Support	Stakeholder-oriented recommendation delivery		Contextual explanation presentation	Trustworthy engineering decisions

The architectural decomposition presented in Table 1 demonstrates that explainability functions as a distributed capability rather than an isolated computational stage. Information continuity across successive modules prevents semantic degradation between engineering observations and final recommendations, strengthening the logical integrity of system reasoning. Such structural integration also improves architectural extensibility because individual modules can evolve independently without disrupting the explanatory chain. Recent reviews similarly emphasize modular explainability architectures as an effective strategy for maintaining scalability while preserving transparent decision support functionality (Ravi et al., 2025).

The decision support layer performs a translational rather than computational function because it converts algorithmic explanations into engineering knowledge that can be interpreted by decision-makers possessing different levels of technical expertise. Effective explanations depend not only on computational correctness but also on their alignment with cognitive expectations and engineering workflows. Stakeholder-oriented explanation design enables identical predictive outputs to be communicated differently while preserving analytical validity. Tailored explainability methodologies proposed for intelligent energy management demonstrate that explanation personalization substantially improves practical decision acceptance without compromising technical rigor (Panagoulas et al., 2023).

Scenario-based architectural analysis further indicates that information traceability remains preserved across engineering applications characterized by varying operational complexity. Manufacturing optimization, infrastructure maintenance, cyber-physical security, and predictive asset management require different analytical priorities, yet the proposed architecture maintains identical information flow principles across these environments. Modular separation between inference and explanation prevents domain-specific customization from weakening conceptual consistency within the overall decision support framework. Comparable observations have emerged from predictive maintenance research where explainability strengthens engineering reasoning despite substantial differences in operational conditions and maintenance objectives (Rajaoarisoa et al., 2025).

Analytical verification also reveals that architectural consistency depends upon explicit dependency relationships among computational reasoning, explanatory generation, and stakeholder interpretation rather than algorithmic sophistication alone. Black-box prediction systems frequently produce technically accurate recommendations whose practical usefulness declines because engineering professionals cannot validate underlying causal relationships. Structured explainability addresses this limitation by preserving logical continuity from engineering evidence to recommended action, reducing ambiguity during high-impact operational decisions. Investigations involving project optimization and engineering asset management similarly identify transparent reasoning as an essential

prerequisite for dependable AI-assisted decision support rather than a supplementary visualization feature (Bisaria et al., 2025; Takon, 2025).

The conceptual architecture developed through this analytical design establishes a coherent foundation for integrating predictive intelligence, explainability mechanisms, and engineering decision support within a unified framework. Interoperability among architectural modules enhances adaptability across diverse engineering domains while preserving transparency, traceability, and human interpretability throughout the decision lifecycle. The proposed framework contributes a systematic architectural perspective that extends existing explainability research beyond algorithm-centered analysis toward comprehensive engineering system design. Emerging discussions concerning real-time explainable decision support likewise indicate that future intelligent engineering systems will increasingly depend upon architectures capable of balancing computational performance with transparent reasoning and accountable decision-making (Priya & Singh, 2026).

Analytical Evaluation of Explainability Mechanisms in Engineering Decision Support Systems

The analytical evaluation demonstrates that explainability should be interpreted as a multidimensional capability that strengthens engineering decision support through the interaction of transparency, interpretability, traceability, and contextual reasoning rather than through a single explanatory technique. Conventional evaluation frameworks frequently emphasize predictive performance while assigning secondary importance to explanation quality, creating an imbalance between algorithmic efficiency and engineering usability. The proposed conceptual framework addresses this imbalance by positioning explanation fidelity as an integral component of system functionality rather than as a supplementary output generated after prediction. Contemporary investigations increasingly recognize that engineering decision support requires simultaneous optimization of predictive reliability and explanation quality to ensure sustainable human–AI collaboration (Kostopoulos et al., 2024; Hussain et al., 2021).

The scenario-driven analytical assessment indicates that explanation fidelity depends on preserving logical consistency between engineering evidence, computational inference, and explanatory outcomes across every architectural layer. Explanations that accurately describe model behavior but fail to reflect engineering semantics contribute little to operational decision-making because practitioners evaluate recommendations within domain-specific contexts instead of purely computational environments. The proposed framework resolves this issue by maintaining explicit dependencies among engineering variables, inference pathways, and explanation generation throughout the decision lifecycle. Similar observations have been reported in Industry 4.0 environments where explainability becomes meaningful only when computational transparency corresponds with operational engineering knowledge (Ahmed et al., 2022).

Interpretability emerges as a distinct analytical dimension because engineering stakeholders possess heterogeneous expertise, professional responsibilities, and decision objectives that influence their understanding of intelligent recommendations. Uniform explanation strategies frequently overlook these cognitive differences, reducing the practical value of otherwise accurate predictive systems. The proposed architecture incorporates adaptive explanation structures capable of communicating identical predictive outcomes through multiple reasoning perspectives while preserving semantic consistency. Investigations concerning supply chain decision support similarly conclude that explainability capabilities must accommodate organizational diversity if intelligent systems are expected to improve collaborative engineering decisions (Olan et al., 2025; Sadeghi et al., 2024).

Transparency also extends beyond algorithm disclosure because engineering decision support requires visibility into how evidence evolves into recommendations through sequential computational reasoning. Exposing mathematical operations alone rarely enables practitioners to validate engineering logic when complex operational dependencies remain implicit. Transparent information flow therefore becomes more valuable than isolated model inspection because it preserves causal continuity between engineering observations and recommended actions.

Comparable analytical perspectives have been proposed for transparent AI decision-making, emphasizing structured reasoning pathways instead of isolated visualization techniques (Praveenraj et al., 2023). The analytical relationships among the principal explainability dimensions become more apparent when they are evaluated collectively within the proposed engineering framework. Table 2 summarizes the conceptual assessment criteria developed through the architecture-based analysis.

Table 2. Analytical Dimensions for Evaluating Explainability in Engineering Decision Support

Evaluation Dimension	Analytical Focus		Engineering Function		Expected Impact	Decision
Explanation Fidelity	Consistency between inference and explanation		Preserves accuracy	reasoning	Higher reliability	decision
Interpretability	Human comprehension of explanations		Supports understanding	engineering	Improved acceptance	stakeholder
Transparency	Visibility of reasoning	computational	Enables validation	engineering	Greater accountability	operational
Traceability	Information continuity across modules		Facilitates auditing	engineering	Better governance and compliance	
Scalability	Adaptability across engineering domains		Supports deployment	modular	Increased implementation flexibility	

The analytical framework summarized in Table 2 illustrates that explainability quality cannot be represented by a single evaluation criterion because engineering decision support simultaneously requires technical, cognitive, and organizational consistency. Mutual interaction among these dimensions determines whether explanations remain operationally meaningful under varying engineering conditions. Architectural evaluation therefore prioritizes balanced performance across multiple explainability characteristics instead of maximizing one criterion at the expense of others. Similar multidimensional assessment principles have been emphasized in engineering-oriented reviews of explainable decision support systems (Ravi et al., 2025; Petrauskas et al., 2023).

Traceability represents another essential analytical outcome because engineering environments increasingly require auditable decision pathways that satisfy technical governance and regulatory accountability. Recommendations lacking traceable reasoning frequently complicate root-cause investigation when unexpected operational consequences emerge after deployment. The proposed framework maintains explicit information continuity between engineering inputs, inference mechanisms, explanatory generation, and stakeholder recommendations, reducing ambiguity during post-decision evaluation. Comparable architectural priorities have been identified in cybersecurity-oriented engineering assets where explainable reasoning improves accountability for high-risk operational decisions (Takon, 2025).

Scalability analysis further indicates that explainability mechanisms should remain independent of specific engineering applications while preserving compatibility with heterogeneous operational contexts. Domain-specific optimization frequently improves localized performance but limits architectural transferability across engineering disciplines characterized by different data structures and decision objectives. Modular explainability enables adaptation without altering the conceptual integrity of the decision support architecture because functional responsibilities remain clearly separated. Studies involving predictive maintenance and real-time engineering project optimization similarly demonstrate that modular explainability architectures facilitate broader implementation while maintaining analytical consistency across operational environments (Rajaoarisoa et al., 2025; Bisaria et al., 2025).

Human-centered evaluation strengthens the conceptual framework by recognizing that engineering decisions ultimately depend on practitioner confidence rather than computational capability alone. Trust develops when explanations remain understandable, logically coherent, and technically relevant throughout engineering workflows instead of merely satisfying algorithmic interpretability metrics. The analytical model therefore positions stakeholder comprehension as an evaluation objective equivalent to computational transparency because effective engineering decision support requires continuous interaction between intelligent systems and human expertise. Comparable findings have been reported across medical decision support environments where explainability substantially improves confidence in AI-assisted recommendations despite increasing algorithmic complexity (Knapič et al., 2021; Xu et al., 2023; Amann et al., 2022).

The analytical evaluation confirms that explainability functions as an engineering quality attribute governing the interaction among predictive intelligence, transparent reasoning, and accountable decision support rather than as an isolated interpretability technique. Balanced integration of fidelity, interpretability, transparency, traceability, and scalability establishes a stronger conceptual foundation for trustworthy engineering intelligence than architectures emphasizing predictive optimization alone. Future Engineering Decision Support Systems are expected to derive competitive advantages from explanation-aware architectures capable of supporting complex operational reasoning while maintaining human oversight throughout the decision process. Recent developments in explainable AI for industrial, healthcare, and engineering applications consistently reinforce this architectural direction as intelligent systems continue expanding into safety-critical and real-time operational environments (Alabdulhafith et al., 2023; Chadaga et al., 2023; Rongali, 2023; Shams et al., 2024).

Theoretical Implications and Future Engineering Directions for Explainable Artificial Intelligence-Based Decision Support Systems

The conceptual analysis indicates that the proposed Explainable Artificial Intelligence framework extends engineering decision support beyond conventional prediction-oriented architectures by incorporating transparent reasoning as an intrinsic system capability. Engineering intelligence increasingly depends on the ability to justify computational recommendations through coherent explanatory pathways that align with operational objectives and stakeholder expectations. Decision quality consequently becomes a product of both predictive competence and explanatory validity because engineering practitioners evaluate recommendations through technical reasoning rather than numerical outputs alone. Recent systematic reviews similarly argue that explainability represents a foundational requirement for the next generation of intelligent decision support systems operating in high-consequence environments (Sahoh & Choksuriwong, 2023; Kostopoulos et al., 2024).

The proposed framework also contributes theoretically by redefining explainability as an architectural property distributed across multiple functional layers instead of limiting it to post-hoc interpretation modules. This analytical perspective establishes explicit dependencies among engineering data representation, computational inference, explanation generation, and stakeholder-oriented decision support, creating a coherent reasoning structure throughout the system lifecycle. Functional integration reduces conceptual fragmentation that frequently characterizes explainability research, where individual algorithms are evaluated independently from broader engineering system design. Comparable perspectives have emerged in recent discussions describing XAI-enabled decision support as an ecosystem requiring coordinated interaction between artificial intelligence and engineering knowledge rather than isolated model optimization (Ravi et al., 2025; Petrauskas et al., 2023).

Engineering governance represents another theoretical dimension strengthened by the proposed architecture because explainability directly influences accountability, verification, and responsible deployment of intelligent systems. Transparent reasoning enables engineering organizations to document the logical progression from operational evidence to computational recommendation, facilitating technical audits and post-decision evaluation. Governance mechanisms become increasingly significant as engineering infrastructures adopt autonomous and semi-autonomous decision support technologies capable of influencing safety, productivity, and resource allocation. Similar observations have been emphasized in engineering studies demonstrating that explainability supports organizational confidence by strengthening accountability within complex cyber-physical environments (Takon, 2025; Ahmed et al., 2022).

The conceptual framework further demonstrates that explanation quality should be evaluated according to engineering relevance rather than computational sophistication alone. Complex explanation algorithms frequently produce technically accurate interpretations while remaining difficult for engineering practitioners to translate into operational decisions. Practical explainability therefore depends upon contextual alignment between computational reasoning and engineering knowledge instead of algorithmic complexity. Recent investigations into transparent AI decision-making similarly conclude that explanation usefulness is determined by contextual applicability rather than mathematical elegance (Praveenraj et al., 2023).

The implications of the proposed framework become more explicit when its theoretical contributions are organized according to engineering design objectives and future implementation priorities. Table 3 summarizes the principal conceptual implications identified through the present analytical evaluation.

Table 3. Theoretical Contributions and Future Engineering Directions

Conceptual Dimension	Analytical Contribution	Engineering Implication	Future Direction	Research
Architectural Explainability	Integrated reasoning across functional modules	Transparent lifecycle	decision	Domain-specific implementation
Human-Centered Intelligence	Stakeholder-oriented explanation strategies	Improved collaboration	engineering	Adaptive explanation models
Engineering Governance	Traceable reasoning	computational Accountable deployment	AI	Explainability standards
Modular Scalability	Technology-independent architecture	Cross-domain interoperability		Large-scale engineering integration
Evaluation Framework	Multi-dimensional explainability assessment	Reliable validation	engineering	Quantitative benchmarking

The conceptual synthesis presented in Table 3 illustrates that explainability influences engineering system performance through interconnected technical and organizational dimensions rather than isolated computational improvements. Future implementation efforts can therefore prioritize architectural refinement without compromising interoperability among heterogeneous engineering applications. Modular development pathways also facilitate incremental integration of emerging explainability techniques as artificial intelligence technologies continue evolving. Similar development trajectories have been proposed for engineering-oriented XAI frameworks emphasizing long-term adaptability instead of application-specific optimization (Rajaoarisoa et al., 2025; Bisaria et al., 2025).

Interdisciplinary integration constitutes another important implication because engineering decision support increasingly operates at the intersection of artificial intelligence, systems engineering, human factors, and organizational management. Effective explainability consequently requires conceptual models capable of accommodating multiple disciplinary perspectives without sacrificing analytical consistency. The proposed architecture provides a common structural foundation through which engineering specialists, data scientists, and decision-makers can interpret computational recommendations using compatible reasoning principles. Comparable interdisciplinary requirements have been identified in supply chain resilience studies where explainable intelligence improves coordination across organizational functions involved in complex decision processes (Sadeghi et al., 2024; Olan et al., 2025).

The analytical framework also highlights opportunities for extending explainability into real-time engineering environments characterized by continuous data streams and adaptive operational conditions. Static explanation mechanisms may become insufficient when engineering systems modify predictive behavior dynamically in response to evolving operational information. Future explainability architectures should therefore incorporate adaptive reasoning capabilities capable of preserving transparency despite continuous model evolution. Emerging investigations concerning real-time explainable decision support identify adaptive explanation generation as an essential research direction for intelligent engineering infrastructures (Priya & Singh, 2026).

Cross-domain applicability further strengthens the theoretical value of the proposed architecture because identical explainability principles remain relevant despite substantial differences among engineering sectors. Manufacturing optimization, predictive maintenance, healthcare engineering, agriculture, construction management, and industrial automation share common requirements for trustworthy computational reasoning despite possessing different operational objectives. The modular conceptual design facilitates knowledge transfer across these domains while preserving engineering-specific customization through independent functional components. Similar patterns have been

documented across healthcare, agriculture, and industrial applications where explainable artificial intelligence consistently improves decision credibility under heterogeneous operational conditions (Khanna et al., 2023; Alabdulhafith et al., 2023; Chadaga et al., 2023; Shams et al., 2024).

The proposed conceptual framework establishes a comprehensive theoretical foundation for advancing Explainable Artificial Intelligence as a core architectural principle within Engineering Decision Support Systems rather than as an auxiliary interpretability feature. Distributed explainability, human-centered reasoning, modular scalability, and engineering governance collectively strengthen the capacity of intelligent systems to support transparent, accountable, and technically reliable decision-making across diverse engineering environments. Future investigations may transform this conceptual architecture into empirical implementations by integrating quantitative performance validation, domain-specific datasets, and standardized explainability benchmarks while preserving the analytical relationships developed in the present framework. Continued progress in explainable engineering intelligence is expected to depend on systematic integration between computational innovation and engineering reasoning capable of sustaining trust, operational robustness, and responsible artificial intelligence deployment (Panagoulas et al., 2023; Zhan et al., 2024; Mahbooba et al., 2021).

CONCLUSION

The present study develops a conceptual Explainable Artificial Intelligence (XAI) framework for Engineering Decision Support Systems by integrating engineering data processing, AI inference, explainability mechanisms, and human-centered decision support into a unified multilayer architecture. The analytical evaluation demonstrates that explanation fidelity, interpretability, transparency, traceability, and modular scalability operate as interdependent engineering quality attributes that determine the effectiveness of intelligent decision support beyond predictive performance alone. Structured integration among these architectural components preserves logical consistency between engineering evidence, computational reasoning, and stakeholder-oriented recommendations while strengthening accountability, governance, and interdisciplinary interoperability across heterogeneous engineering domains. The proposed framework also repositions explainability from a post-hoc interpretability technique to an intrinsic architectural capability that supports trustworthy engineering intelligence through transparent reasoning pathways. The study contributes theoretically by extending existing XAI discourse toward an architecture-centered engineering perspective and contributes methodologically by providing a reproducible analytical framework suitable for future empirical validation, quantitative benchmarking, and domain-specific implementation in intelligent engineering systems operating under complex, dynamic, and safety-critical conditions.

REFERENCES

- Ahmed, I., Jeon, G., & Piccialli, F. (2022). From artificial intelligence to explainable artificial intelligence in industry 4.0: a survey on what, how, and where. *IEEE transactions on industrial informatics*, 18(8), 5031-5042.
- Alabdulhafith, M., Saleh, H., Elmannai, H., Ali, Z. H., El-Sappagh, S., Hu, J. W., & El-Rashidy, N. (2023). A clinical decision support system for edge/cloud ICU readmission model based on particle swarm optimization, ensemble machine learning, and explainable artificial intelligence. *IEEE Access*, 11, 100604-100621.
- Amann, J., Vetter, D., Blomberg, S. N., Christensen, H. C., Coffee, M., Gerke, S., ... & Z-Inspection Initiative. (2022). To explain or not to explain?—Artificial intelligence explainability in clinical decision support systems. *PLOS Digital Health*, 1(2), e0000016. <https://doi.org/10.1371/journal.pdig.0000016>.
- Bisaria, C., Cherian, E., Abass, S., Jumaa, A. H., & Miys, H. S. (2025). Explainable AI-Based Decision Support System for Real-Time Project Management Optimization. In *2025 International Conference on Recent Innovation in Science Engineering and Technology (ICRISET)* (pp. 1-7). IEEE.
- Chadaga, K., Prabhu, S., Bhat, V., Sampathila, N., Umakanth, S., & Chadaga, R. (2023). A decision support system for diagnosis of COVID-19 from non-COVID-19 influenza-like illness using explainable artificial intelligence. *Bioengineering*, 10(4), 439. <https://doi.org/10.3390/bioengineering10040439>.

- Cochran, D. S., Smith, J., Mark, B. G., & Rauch, E. (2022, June). Information model to advance explainable AI-based decision support systems in manufacturing system design. In *International Symposium on Industrial Engineering and Automation* (pp. 49-60). Cham: Springer International Publishing. https://doi.org/10.1007/978-3-031-14317-5_5.
- Hussain, F., Hussain, R., & Hossain, E. (2021). Explainable artificial intelligence (XAI): An engineering perspective. *arXiv preprint arXiv:2101.03613*. <https://doi.org/10.48550/arXiv.2101.03613>.
- Khanna, V. V., Chadaga, K., Sampathila, N., Chadaga, R., Prabhu, S., KS, S., ... & Bhat, D. (2023). A decision support system for osteoporosis risk prediction using machine learning and explainable artificial intelligence. *Heliyon*, 9(12).
- Knapič, S., Malhi, A., Saluja, R., & Främling, K. (2021). Explainable artificial intelligence for human decision support system in the medical domain. *Machine Learning and Knowledge Extraction*, 3(3), 740-770. <https://doi.org/10.3390/make3030037>.
- Kostopoulos, G., Davrazos, G., & Kotsiantis, S. (2024). Explainable artificial intelligence-based decision support systems: A recent review. *Electronics*, 13(14), 2842. <https://doi.org/10.3390/electronics13142842>.
- Mahbooba, B., Timilsina, M., Sahal, R., & Serrano, M. (2021). Explainable artificial intelligence (XAI) to enhance trust management in intrusion detection systems using decision tree model. *Complexity*, 2021(1), 6634811. <https://doi.org/10.1155/2021/6634811>.
- Olan, F., Spanaki, K., Ahmed, W., & Zhao, G. (2025). Enabling explainable artificial intelligence capabilities in supply chain decision support making. *Production Planning & Control*, 36(6), 808-819. <https://doi.org/10.1080/09537287.2024.2313514>.
- Panagoulas, D. P., Sarmas, E., Marinakis, V., Virvou, M., Tsihrintzis, G. A., & Doukas, H. (2023). Intelligent decision support for energy management: A methodology for tailored explainability of artificial intelligence analytics. *Electronics*, 12(21), 4430. <https://doi.org/10.3390/electronics12214430>.
- Petrauskas, V., Jasinevicius, R., Kazanavicius, E., & Meskauskas, Z. (2023). The paradigm of an explainable artificial intelligence (XAI) and data science (DS)-based decision support system (DSS). In *Data Science in Applications* (pp. 167-209). Cham: Springer International Publishing.
- Praveenraj, D. D. W., Victor, M., Vennila, C., Hussein Alawadi, A., Diyora, P., Vasudevan, N., & Avudaiappan, T. (2023). Exploring explainable artificial intelligence for transparent decision making. In *E3S Web of Conferences* (Vol. 399, p. 04030). EDP Sciences. <https://doi.org/10.1051/e3sconf/202339904030>.
- Priya, B., & Singh, P. (2026). Explainable AI Models for Real-Time Decision Support Systems. *International Journal of Engineering & Extended Technologies Research (IJEETR)*, 8(1), 138-145. <https://doi.org/10.15662/IJEETR.2026.0801016>.
- Rajaoarisoa, L., Randrianandraina, R., Nalepa, G. J., & Gama, J. (2025). Decision-making systems improvement based on explainable artificial intelligence approaches for predictive maintenance. *Engineering Applications of Artificial Intelligence*, 139, 109601. <https://doi.org/10.1016/j.engappai.2024.109601>.
- Ravi, M., Negi, A., Bommi, N. S., & Rouf, N. (2025). Evolution of AI-driven decision making with decision support systems, expert systems, recommender systems, and XAI. *IETE Technical Review*, 42(4), 428-465. <https://doi.org/10.1080/02564602.2025.2512086>.
- Rongali, S. K. (2023). Explainable Artificial Intelligence (XAI) Framework for Transparent Clinical Decision Support Systems. *International Journal of Medical Toxicology and Legal Medicine*, 26(3), 22-31.
- Sadeghi, K., Ojha, D., Kaur, P., Mahto, R. V., & Dhir, A. (2024). Explainable artificial intelligence and agile decision-making in supply chain cyber resilience. *Decision Support Systems*, 180, 114194. <https://doi.org/10.1016/j.dss.2024.114194>.
- Sahoh, B., & Choksuriwong, A. (2023). The role of explainable Artificial Intelligence in high-stakes decision-making systems: a systematic review. *Journal of Ambient Intelligence and Humanized Computing*, 14(6), 7827-7843.
- Shams, M. Y., Gamel, S. A., & Talaat, F. M. (2024). Enhancing crop recommendation systems with explainable artificial intelligence: a study on agricultural decision-making. *Neural Computing and Applications*, 36(11), 5695-5714.

- Takon, A. (2025). Explainable AI for Threat Modelling and Decision Support in Engineering Assets. *Journal of Cyber-Physical Security and Robotics*, 1(02), 46-52.
- Xu, Q., Xie, W., Liao, B., Hu, C., Qin, L., Yang, Z., ... & Luo, A. (2023). Interpretability of clinical decision support systems based on artificial intelligence from technological and medical perspective: a systematic review. *Journal of healthcare engineering*, 2023(1), 9919269. <https://doi.org/10.1155/2023/9919269>.
- Zhan, J., Fang, W., Love, P. E., & Luo, H. (2024). Explainable artificial intelligence: Counterfactual explanations for risk-based decision-making in construction. *IEEE Transactions on Engineering Management*, 71, 10667-10685.