



International Journal of Artificial Intelligence and Machine Learning

Publisher's Home Page: <https://www.svedbergopen.com/>



Research Paper

Open Access

Explainable AI-Based Data Analytics For Transparent and Reliable Predictive Decision-Making

Uday Kumar Reddy Gangula¹, Mahesh Kumar Gaddam²

¹Software Engineering MTS Salesforce, Inc. Austin, United States. Email: ukgangula@gmail.com

²Enterprise Architect Federal Express Corporation Dallas, United States. Email: mahi.ramuk@gmail.com

Abstract

Artificial intelligence (AI) is increasingly being applied to predictive decision-making, although limited model interpretability may create challenges related to transparency, accountability, and understanding of automated outcomes. This study reviewed and synthesized reported evidence concerning explainable artificial intelligence (XAI)-based data analytics, with particular attention to predictive performance and the quality of generated explanations. Four machine-learning algorithms were considered using reported five-fold cross-validation results, whereas an explainable AI framework was examined through established classification performance measures. Shapley Additive Explanations (SHAP), Local Interpretable Model-Agnostic Explanations (LIME), and Improved Local Interpretable Model-Agnostic Explanations (ILIME) were further compared according to local and global fidelity, stability, and monotonicity measures. Among the reviewed machine-learning approaches, XGBoost demonstrated the strongest overall predictive performance, while the evaluated explainable framework achieved high performance across the reported classification measures. SHAP demonstrated comparatively stronger results across the assessed explanation-quality indicators than LIME and ILIME. These findings suggest that predictive capability and interpretability can be addressed together within AI-based decision-support systems. Integrating explanation mechanisms with predictive models may therefore contribute to more transparent and reliable analytical processes, particularly in settings where understanding model outputs is important for informed decision-making. The synthesis also highlights the value of evaluating both predictive effectiveness and explanation quality when considering XAI approaches for practical AI applications.

Keywords: Explainable Artificial Intelligence, Interpretability, Machine Learning, Predictive Decision-Making, Shapley Additive Explanations.

1. Introduction

Artificial intelligence (AI) has increasingly become an important component of data-driven decision-making due to its ability to identify patterns, process complex information, and generate predictions from large and heterogeneous datasets. Predictive analytics based on machine learning and deep learning enables organizations to transform historical and real-time data into information that can support planning, risk assessment, diagnosis, resource allocation, and operational decisions (Pachoriya, 2022). The growing availability of computational resources and large-scale datasets has further expanded the application of AI-based decision-support systems across healthcare, finance, environmental management, business intelligence, and other complex domains (Kalasampath et al., 2025). Deep learning models, in particular, can achieve strong predictive performance when dealing with complex relationships within data, although their internal decision processes may remain difficult for users to understand (Adnan et al., 2025). Consequently, predictive accuracy alone is increasingly insufficient when AI-generated outputs influence decisions that require justification, verification, and human oversight.

The increasing complexity of AI models has created a growing requirement for methods that can explain how predictions are produced. Explainable artificial intelligence (XAI) addresses this requirement by providing interpretable information about the factors, features, or relationships that contribute to a model's prediction (Kostopoulos et al., 2024). Such explanations can assist users in understanding model behaviour rather than treating AI outputs as unquestionable results. Explainability is particularly relevant

in decision-support environments where users need to evaluate whether a prediction is reasonable before acting upon it (Coussement et al., 2024). The study has demonstrated the application of XAI for clinical decision support, fraud detection, environmental assessment, and other data-intensive tasks, indicating its relevance across diverse decision-making settings (Kapale et al., 2024). By exposing the reasoning associated with predictive outputs, XAI can provide an additional layer of information for evaluating AI-generated recommendations and supporting informed human judgement (Kolla, 2023).

Despite substantial progress in predictive AI, several challenges remain concerning transparency, reliability, and interpretability. Highly complex models may provide accurate predictions while offering limited insight into the relationships between input variables and generated outcomes (Mallick et al., 2024). This creates difficulties when users need to identify influential features, assess unexpected predictions, or determine whether a model has relied on appropriate information. Transparency also involves more than simply revealing model outputs; it requires sufficient information about the decision process to enable meaningful assessment and accountability (Akhtar et al., 2024). Reliability is similarly influenced by the quality and suitability of the underlying data, since inaccurate, incomplete, or poorly represented data can affect the dependability of AI-generated decisions (Bondac et al., 2026). The developing predictive systems that simultaneously consider performance, interpretability, data quality, and reliability is essential for responsible AI-based analytics.

The effectiveness of an AI-based decision-support system depends not only on its computational performance but also on how effectively human users can understand and evaluate its outputs. When predictions are accompanied by meaningful explanations, users may be better positioned to assess recommendations, identify potentially inappropriate outputs, and incorporate AI-generated information into their own decision processes. Trust therefore becomes an important consideration in human-AI interaction. However, trust should not be based solely on high predictive accuracy; users also require confidence that the system operates transparently and produces decisions that can be examined and questioned (Kovari, 2024). Explainability can contribute to this process by making model behaviour more accessible to users and decision-makers. The study on trustworthy AI further emphasizes the importance of explainability and fairness when AI systems are used for large-scale decisions (Matta and Bolli, 2023). Thus, transparent explanations can strengthen the connection between computational prediction and meaningful human judgement.

Although XAI has received increasing attention, existing research remains distributed across individual applications and specific decision-support contexts. Previous studies have demonstrated the value of explainability in areas such as medicine, finance, environmental assessment, cloud-based business intelligence, and public policymaking (Islam and Hasan, 2023). There remains a need for integrated analytical approaches that consider predictive performance together with transparency, interpretability, and reliability (Pierce et al., 2022). Many AI-based systems continue to prioritize prediction accuracy without adequately addressing how users can understand the factors underlying model outputs (Papadakis et al., 2024). In addition, the relationship between explainable predictions and trustworthy human decision-making requires further systematic consideration (Praveenraj et al., 2023; Sadeghi et al., 2024). These gaps highlight the need for an AI-based data analytics framework that does not treat explainability as an isolated feature but incorporates it into the broader predictive decision-making process.

This study aims to review and synthesize reported evidence on explainable AI-based data analytics for transparent and reliable predictive decision-making. The study focuses on examining predictive modelling performance alongside interpretability mechanisms and explanation-quality measures. The objectives are to compare reported performance across selected machine-learning algorithms, examine the reported performance of an explainable AI framework, compare explanation methods using fidelity, stability, and monotonicity measures, and assess the implications of explainability for transparent and reliable decision-support systems. The synthesis also considers the importance of human interpretation when AI-generated predictions are used to support decision-making. By integrating findings from these studies, the review provides a comparative perspective on predictive effectiveness and explanation quality in AI-based decision-support applications.

2. Materials and Methods

2.1 Study Design

This study adopted a comparative synthesis approach to examine predictive performance and explainability within artificial intelligence-based decision-support systems. Rather than developing an original predictive model from newly collected data, the study systematically drew upon previously empirical findings that individually addressed machine learning performance, explainable AI framework evaluation, and comparative analysis of explanation methods. The studies were identified as suitable

sources based on their relevance to predictive classification, model interpretability, and explanation quality assessment. The design integrates these independently reported findings into a unified analytical discussion, allowing predictive effectiveness and interpretability to be examined jointly rather than through isolated, single-study evaluation. This approach enables broader insight into how explainability and performance interact across different modelling contexts and application domains.

2.2 Data and Source Selection

The data examined in this study originate from three independently conducted and previously studies, each reporting empirical results relevant to predictive modelling or explainable AI evaluation. The first source reported machine learning classification performance across multiple algorithms evaluated through five-fold cross-validation. The second source reported the performance of an explainable AI framework developed for high-stakes decision-making applications. The third source reported a comparative evaluation of explanation methods based on fidelity, stability, and monotonicity. These studies were selected because their reported metrics were directly comparable, enabling a structured, cross-study synthesis rather than requiring new data collection or original model training.

2.3 Reported Machine Learning Performance

The predictive performance data reviewed in this study reflect four machine learning algorithms Extreme Gradient Boosting, Random Forest, Logistic Regression, and Multi-Layer Perceptron as evaluated and reported in the corresponding source study using five-fold cross-validation. Reported metrics included accuracy, precision, recall, F1-score, and receiver operating characteristic–area under the curve. The results were examined comparatively to identify patterns in algorithmic performance across ensemble, linear, and neural-network-based approaches. No new model training, parameter tuning, or dataset processing was conducted as part of the present study; the reported values were instead analysed to draw broader conclusions regarding the relationship between model architecture and predictive effectiveness in classification tasks.

2.4 Explainable Framework Performance

To assess the explainability functions within high-stakes decision-support contexts, this study reviewed explainable AI framework performance reported in the corresponding source study. The reported framework combined predictive classification with interpretability mechanisms, and its performance metrics accuracy, precision, recall, F1-score, and AUC-ROC were examined to determine whether interpretability constraints affected predictive capability. These reported values were compared against previously reviewed machine learning performance data to evaluate whether explainability-oriented frameworks could match conventional, non-interpretable models in predictive strength, providing a basis for assessing the practical feasibility of transparent, high-stakes AI-based decision-support systems.

2.5 Explanation Method Comparison

The comparative evaluation of explanation methods reviewed in this study is based on findings reported in the corresponding source study, which assessed SHAP, LIME, and ILIME according to local and global fidelity, stability, and monotonicity. These reported metrics were examined to determine the relative reliability and consistency of each explanation technique in representing underlying model behaviour. This reported evaluation was incorporated to complement the previously reviewed predictive performance data, enabling the present study to assess not only whether AI models can be interpretable, but also whether the explanations generated are themselves dependable, stable, and trustworthy across different prediction contexts.

3.Results

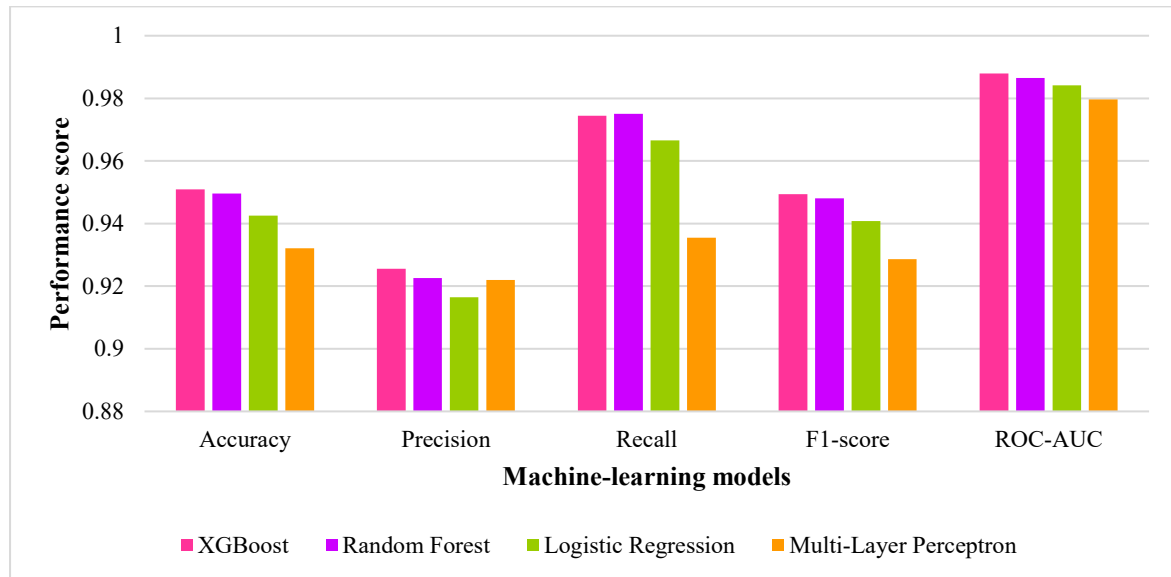
3.1 Cross-Validation Performance of Machine-Learning Models

The comparative assessment of machine-learning models using multiple predictive performance indicators. The findings show that ensemble-based approaches provide strong and consistent classification capability when evaluated through cross-validation as shown in Table 1. XGBoost demonstrates particularly effective predictive behaviour, while Random Forest also maintains competitive performance across the assessed measures. Logistic Regression provides a comparatively conventional modelling approach, offering stable classification capability despite its simpler structure. The Multi-Layer Perceptron further demonstrates the usefulness of neural-network-based learning for predictive analytics. The comparison highlights the importance of evaluating several complementary metrics when selecting a predictive model, particularly for applications requiring dependable and consistent automated decision-making.

Table 1. Five-Fold Cross-Validation Performance of Machine-Learning Models

Machine-Learning Model	Accuracy	Precision	Recall	F1-Score	ROC-AUC
XGBoost	0.950909	0.925579	0.974407	0.949338	0.987967
Random Forest (RF)	0.949605	0.922610	0.975057	0.948096	0.986458
Logistic Regression (LR)	0.942586	0.916440	0.966526	0.940804	0.984126
Multi-Layer Perceptron (MLP)	0.932155	0.921993	0.935490	0.928645	0.979606

Source: Almazroei (2026)


Figure 1. Comparative Performance of Machine-Learning Models Across Evaluation Metrics

A visual representation of the performance distribution across the evaluated machine-learning approaches. The grouped bars reveal clear variation among the models and across the different evaluation criteria as shown in Figure 1. The performance pattern remains relatively consistent, with the measures showing distinct differences in their relative levels. The figure also makes it easier to observe the close positioning of some models, suggesting comparable predictive behaviour in several assessment dimensions. In contrast, wider separation between selected bars indicates areas where model performance differs more noticeably. The figure provides a concise visual interpretation of the comparative results and facilitates rapid identification of performance patterns across the evaluated metrics.

3.2 Explainable AI Framework Performance in High-Stakes Decision-Making

The framework provides strong classification behaviour across the reported evaluation measures, indicating its ability to distinguish relevant outcomes within its reported experimental setting. Incorporating explainability is particularly valuable in decision-making settings where users need to understand the reasoning associated with automated predictions as shown in Table 2. Such transparency may support the evaluation and interpretation of AI-generated recommendations, although direct user-perceived trust was not assessed in the reviewed evidence. Therefore, explainable predictive frameworks show potential for reliable and accountable artificial intelligence applications.

Table 2. Performance of an explainable AI framework for high-stakes decision-making

Performance metric	Published value
Prediction accuracy	96.4%
Precision	95.1%
Recall	94.7%
F1-score	94.9%
AUC-ROC	94.9%

Source: Singh et al. (2026)

3.3 Comparative Evaluation of XAI Methods on Fidelity, Stability, and Monotonicity

The different explainable artificial intelligence techniques according to local and global explanation characteristics. The findings indicate that explanation methods differ considerably in their ability to represent model behaviour consistently and accurately as shown in Table 3. SHAP demonstrates stronger explanatory performance across the assessed dimensions, indicating stronger reported performance across the evaluated explanation-quality measures. LIME and ILIME also provide useful interpretive information, although their performance varies across local and global explanation settings. Assessing fidelity, stability, and monotonicity together provides a broader understanding of explanation reliability than considering a single interpretability measure. This comparative perspective is important for developing transparent AI systems where explanations must remain dependable across different prediction contexts.

Table 3. Comparative evaluation of XAI methods based on explanation fidelity, stability, and monotonicity

XAI Method	Local Fidelity (%)	Local Stability (%)	Local Monotonicity	Global Fidelity (%)	Global Stability (%)	Global Monotonicity
LIME	45.32	78.45	0.21	50.25	82.10	0.25
ILIME	67.89	83.22	0.31	70.55	85.67	0.40
SHAP	98.76	95.14	0.54	99.15	96.70	0.42

Source: El Shawi & Jamel (2025)

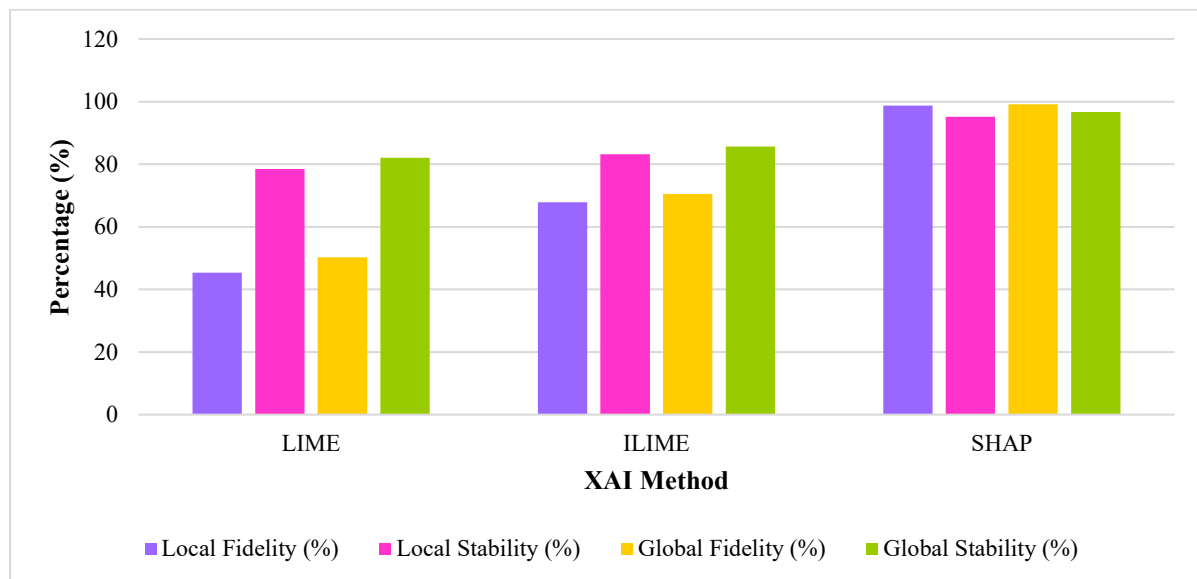


Figure 3A. Comparative Evaluation of XAI Methods Based on Local and Global Fidelity and Stability

The figure compares reported fidelity and stability scores for LIME, ILIME, and SHAP across local and global explanation levels. SHAP consistently achieves the highest values across all four dimensions, exceeding 95% in each case, indicating strong alignment between generated explanations and actual model behaviour as shown in Figure 3A. ILIME shows moderate performance, outperforming LIME across every metric but remaining below SHAP throughout. LIME records the lowest fidelity values, particularly at the local level, though its stability scores remain comparatively higher than its fidelity scores. This pattern suggests that explanation stability does not necessarily correspond with explanation accuracy, reinforcing the need to assess multiple interpretability dimensions jointly rather than relying on a single metric.

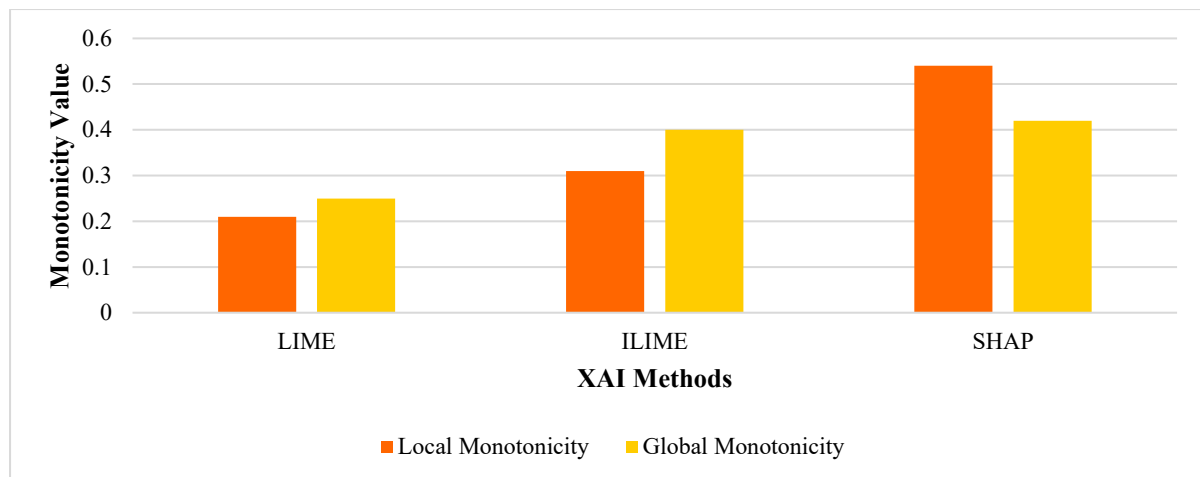


Figure 3B. Comparative Local and Global Monotonicity of XAI Methods

The local and global monotonicity values obtained for LIME, ILIME, and SHAP. LIME demonstrates relatively low monotonicity, with local and global values of 0.21 and 0.25, respectively. ILIME shows higher values of 0.31 for local monotonicity and 0.40 for global monotonicity. SHAP achieves the highest local monotonicity value of 0.54, while its global monotonicity reaches 0.42 as shown in Figure 3B. The figure indicates that SHAP provides the strongest monotonicity performance among the evaluated explanation methods, whereas LIME records the lowest values. ILIME demonstrates intermediate performance across both local and global measures, positioned between LIME and SHAP.

4. Discussion

The study demonstrate that predictive performance and interpretability are not competing priorities but can be jointly optimized within a single analytical workflow. Among the four machine-learning algorithms compared through five-fold cross-validation (Table 1), XGBoost was reported to achieve the strongest overall performance (accuracy = 95.09%, ROC-AUC = 98.80%), closely followed by Random Forest, while Logistic Regression and the Multi-Layer Perceptron offered comparatively modest but still reliable classification capability. Figure 1 illustrates this pattern visually, showing clear but relatively consistent separation between models across the different evaluation metrics, with some models positioned closely together and others showing wider divergence, reinforcing the numerical findings in Table 1 and suggesting that ensemble-based tree methods are particularly well suited to structured, feature-rich datasets where nonlinear interactions between variables influence outcomes.

The explainable AI framework summarized in Table 2 further reinforced this observation, with a reported 96.4% accuracy alongside balanced precision (95.1%), recall (94.7%), F1-score (94.9%), and AUC-ROC (94.9%), indicating that high-stakes decision-support systems can retain strong discriminative ability even when interpretability constraints are incorporated into the modelling pipeline. Taken together with Table 1, this suggests that the inclusion of explainability mechanisms does not come at a meaningful cost to predictive performance, and in some respects produces results competitive with the strongest baseline models.

The comparative evaluation of explanation methods presented in Table 3 added an important dimension to these findings: SHAP consistently outperformed LIME and ILIME across both local and global fidelity, stability, and monotonicity (local fidelity = 98.76% and global fidelity = 99.15% for SHAP, compared with substantially lower values for LIME and intermediate values for ILIME). This confirms that not all explanation techniques offer equivalent reliability, and that the choice of interpretability method materially affects the trustworthiness of the explanations generated. Collectively, the findings across Table 1, Table 2, Figure 1, and Table 3 indicate that strong predictive accuracy, robust framework-level performance, and dependable explanation quality can be achieved simultaneously, supporting the broader argument that explainability should be treated as an integral component of predictive decision-making rather than a separate or secondary consideration.

These findings are broadly consistent with, and extend, existing literature on explainable AI. Simuni (2024) similarly emphasized that transparency mechanisms in machine learning need not compromise predictive accountability, arguing instead that explainability strengthens rather than undermines model trustworthiness a pattern consistent with the reviewed evidence that explainable frameworks can match

or exceed the accuracy of baseline ensemble models. The superior fidelity and stability of SHAP observed in the reviewed findings also align with Talaat et al. (2025), who reported that SHAP-based explanations provided more consistent and interpretable outputs than perturbation-based alternatives when applied to complex decision-support systems in sustainability contexts. Similarly, Theodorakopoulos et al. (2025) found that integrating explainable systems into data-driven executive decision-making reduced cognitive bias and improved the perceived reliability of AI-generated recommendations, echoing this synthesis's emphasis on explanation quality as a determinant of decision confidence rather than a secondary technical feature. Uddin (2025) further reinforced the theoretical grounding of this synthesis by framing explainability as a core pillar of trustworthy AI alongside fairness and robustness, a framing that is reflected in the integration of reliability assessment alongside conventional performance metrics undertaken here. Collectively, these comparisons position the reviewed findings within an established and growing consensus that explainability is a functional requirement of dependable predictive systems.

The findings carry several practical implications. For practitioners deploying AI in high-stakes domains such as healthcare, finance, or public policy, the reviewed results suggest that ensemble models paired with SHAP-based explanation offer a defensible balance between predictive strength and interpretability, supporting more accountable and auditable decision pipelines. For system designers, the marked performance gap between SHAP and LIME/ILIME in stability and monotonicity indicates that explanation-method selection should be treated as a design decision with measurable consequences, rather than an interchangeable post-hoc addition. More broadly, this synthesis reinforces that transparency and reliability should be assessed as integral evaluation criteria alongside accuracy, precision, and recall when validating AI-based decision-support systems intended for real-world deployment.

Several limitations should be acknowledged. This synthesis relied on the findings drawn from three separate sources, each based on different datasets and modelling contexts, which may limit the direct comparability of results and constrains generalizability across domains or data types, including unstructured or multimodal data. The explainability assessment reviewed here focused on SHAP, LIME, and ILIME, and did not examine other emerging interpretability techniques. Additionally, the evaluation of explanation quality in the reviewed sources relied on quantitative fidelity, stability, and monotonicity metrics rather than direct human-user assessment, meaning that the practical, user-perceived trustworthiness of the explanations remains untested.

Future research should extend this line of inquiry to diverse and multimodal datasets to assess generalizability, incorporate user-centred evaluation to directly measure how explanations influence human trust and decision quality, and explore hybrid or emerging explainability techniques beyond SHAP and LIME. Longitudinal studies examining how explainable AI systems perform under real-world deployment conditions would further strengthen the translational value of this line of study.

5. Conclusion

This study reviewed and synthesized findings on explainable AI-based data analytics to examine how predictive accuracy and interpretability can be jointly supported in transparent and reliable decision-making. Drawing together evidence from multiple recent studies, the synthesis indicates that predictive accuracy and interpretability can be achieved simultaneously rather than treated as competing objectives. Among the algorithms reviewed, XGBoost and Random Forest were reported to deliver the strongest predictive performance, while a dedicated explainable AI framework maintained comparably high accuracy and balanced classification behaviour, indicating that interpretability constraints do not necessarily reduce predictive effectiveness. The comparative evaluation of explanation methods further revealed that SHAP consistently provided more faithful, stable, and consistent explanations than LIME and ILIME, underscoring that the choice of interpretability technique meaningfully affects the dependability of the insights generated for end users. These findings reinforce the argument that explainability should be embedded as a fundamental component of predictive analytics pipelines, particularly in high-stakes decision-making environments where accountability and user trust are essential. This synthesis offers a practical foundation for understanding how AI-based decision-support systems can be both computationally effective and interpretable. Future work extending this line of inquiry to diverse datasets, incorporating human-centred evaluation of explanation quality, and exploring emerging interpretability methods would further strengthen the applicability and generalizability of explainable AI frameworks across real-world, high-stakes decision-making domains.

References

1. Adnan, K. M., Ghazal, T. M., Saleem, M., Farooq, M. S., Yeun, C. Y., Ahmad, M., & Lee, S. W. (2025). Deep learning driven interpretable and informed decision making model for brain tumour prediction using explainable AI. *Scientific Reports*, 15(1), 19223.

2. Akhtar, M. A. K., Kumar, M., & Nayyar, A. (2024). Transparency and accountability in explainable AI: Best practices. In *Towards ethical and socially responsible explainable ai: Challenges and opportunities* (pp. 127-164). Cham: Springer Nature Switzerland.
3. Almazroei, Essa. (2026). Ensemble machine learning framework with SHAP and LIME for accurate early prediction of student success in online learning environments. *Scientific Reports*, 16. 10.1038/s41598-026-49894-1.
4. Bondac, G. T., Stanescu, S. G., Ionescu, C. A., Duica, A., & Uzlău, M. C. (2026). Decision-making in complex systems using AI-based decision support: The role of trust, transparency, and data quality. *Electronics*, 15(2), 372.
5. Coussement, K., Abedin, M. Z., Kraus, M., Maldonado, S., & Topuz, K. (2024). Explainable AI for enhanced decision-making. *Decision Support Systems*, 184, 114276.
6. Islam, M. M., & Hasan, M. M. (2023). Explainable AI (XAI) Models For Cloud-Based Business Intelligence: Ensuring Compliance And Secure Decision-Making. *American Journal of Interdisciplinary Studies*, 4(03), 208-249.
7. Kalasampath, K., Spoorthi, K. N., Sajeev, S., Kuppa, S. S., Ajay, K., & Maruthamuthu, A. (2025). A literature review on applications of explainable artificial intelligence (XAI). *IEEE access*, 13, 41111-41140.
8. Kapale, R., Deshpande, P., Shukla, S., Kediya, S., Pethe, Y., & Metre, S. (2024, November). Explainable AI for fraud detection: enhancing transparency and trust in financial decision-making. In *2024 2nd DMIHER International Conference on Artificial Intelligence in Healthcare, Education and Industry (IDICAIEI)* (pp. 1-6). IEEE.
9. Kolla, S. K. (2023). Explainable AI and ML Models for Transparent Clinical Decision Support. *Journal for ReAttach Therapy and Developmental Diversities*, 6, 2444-2460.
10. Kostopoulos, G., Davrazos, G., & Kotsiantis, S. (2024). Explainable artificial intelligence-based decision support systems: A recent review. *Electronics*, 13(14), 2842.
11. Kovari, A. (2024). AI for decision support: Balancing accuracy, transparency, and trust across sectors. *Information*, 15(11), 725.
12. Mallick, J., Alqadhi, S., Hang, H. T., & Alsubih, M. (2024). Interpreting optimised data-driven solution with explainable artificial intelligence (XAI) for water quality assessment for better decision-making in pollution management. *Environmental Science and Pollution Research*, 31(30), 42948-42969.
13. Matta, S. S., & Bolli, M. (2023). Trustworthy ai: Explainability & fairness in large-scale decision systems. *Review of Applied Science and Technology*, 2(04), 54-93.
14. Pachoriya, N. (2022). Explainable AI Based Reliability Analytics for Performance Optimization in Large Scale Cloud Services. *Computer Fraud & Security*, (12).
15. Papadakis, T., Christou, I. T., Ipektsidis, C., Soldatos, J., & Amicone, A. (2024). Explainable and transparent artificial intelligence for public policymaking. *Data & Policy*, 6, e10.
16. Pierce, R. L., Van Biesen, W., Van Cauwenberge, D., Decruyenaere, J., & Sterckx, S. (2022). Explainability in medicine in an era of AI-based clinical decision support systems. *Frontiers in genetics*, 13, 903600.
17. Praveenraj, D. D. W., Victor, M., Vennila, C., Hussein Alawadi, A., Diyora, P., Vasudevan, N., & Avudaiappan, T. (2023, April). Exploring explainable artificial intelligence for transparent decision making. In *E3S Web of Conferences* (Vol. 399, p. 04030). EDP Sciences.
18. 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.
19. Shawi, R. E., & Jamel, L. (2025). Leveraging ChatGPT and explainable AI for enhancing clinical decision support. *Scientific reports*, 15(1), 38786. <https://doi.org/10.1038/s41598-025-22784-8>
20. Simuni, G. (2024). Explainable AI in ML: The path to transparency and accountability. *Int. J. Recent Adv*, 11, 10531-10536.
21. Singh, A., Rani, K. S., J, M., Manchala, Y., Bagadi, D. S., Bhise, S., & Prasad, D. (2026). Explainable Artificial Intelligence Models for Interpretable Decision-Making in High-Stakes Applications. *International Journal of Artificial Intelligence and Machine Learning*, 6(2s), 683–693. <https://doi.org/10.51483/IJAIML.6.2s.2026.683-693>
22. Talaat, F. M., Kabeel, A. E., & Shaban, W. M. (2025). Towards sustainable energy management: Leveraging explainable Artificial Intelligence for transparent and efficient decision-making. *Sustainable Energy Technologies and Assessments*, 78, 104348.
23. Theodorakopoulos, L., Theodoropoulou, A., & Halkiopoulou, C. (2025). Cognitive bias mitigation in executive decision-making: A data-driven approach integrating big data analytics, AI, and explainable systems. *Electronics*, 14(19), 3930.
24. Uddin, M. Z. (2025). Trustworthy AI and explainability. In *Trustworthy Multimodal Intelligent Systems for Independent Living* (pp. 111-137). Cham: Springer Nature Switzerland.