

Explainable Machine Learning Pipelines for Customer Risk Scoring in Anti-Money Laundering: A Management and Governance Perspective

Pristly Turjo Mazumder*

Georgia State University, United States.

ABSTRACT

The rising application of machine learning (ML) in the context of the Anti-Money Laundering (AML) systems has improved the ability to identify suspicious activities by customers, but the obscurity of most ML models is a cause to question the issue of transparency, accountability, and regulatory adherence. The study suggests a customer risk scoring explainable machine learning pipeline that incorporates explainable artificial intelligence (XAI) methods and effective management and governance structures. Based on socio-technical and responsible AI governance lenses, the analyzing paper forms a conceptual and empirical framework that consists of a model performance and explainability measurements, applicable to the compliance officers and regulators. The proposed pipeline utilizes interpretable modeling techniques like SHAP and LIME to identify high-risk customers using the AML data that can be observed in real life and give a clear and audited explanation of model decisions. Results indicate that explainable pipelines enhance accuracy in detection as well as build stakeholder trust, justification of decisions, and conformity with emerging regulatory issues like the EU AI Act, and the financial model risk management policies. The research is theoretically and practically important in that it provides a governance-based framework of deploying credible, interpretable ML systems in AML settings, which eventually provides a solution to the discrepancy between interpretability of technical models and managerial responsibility.

Keywords: Explainable AI (XAI), Machine Learning Pipelines, Anti-Money Laundering (AML), Customer Risk Scoring, Governance, Compliance, Responsible AI, Model Transparency.

Journal of Data Analysis and Critical Management (2025)

DOI: 10.64235/dphq6r10

INTRODUCTION

Background and Motivation

Machine Learning (ML) and Artificial Intelligence (AI) are changing the nature of decision-making in highly regulated industries with the promise of revealing complicated trends in large data volumes that could not be discerned by human analysts. Predictive analytics powered by ML has already shown an impressive ability to identify and treat complex illnesses by using information-driven decision support systems in the field of healthcare (Manik et al., 2021; Osamika et al., 2023; Chakilam, 2022). Equally, multimodal AI models have demonstrated potential in using multiple data sources to improve predictive accuracy and understandability (Soenksen et al., 2022; Kline et al., 2022). The innovations highlight the disruptive nature of AI when implemented in the context of sound governance and data management strategies (Acosta et al., 2022; Bardhan et al., 2020).

Corresponding Author: Pristly Turjo Mazumder, Georgia State University, United States, e-mail: pristly.turjo@gmail.com

How to cite this article: Mazumder, P.T. (2025). Explainable Machine Learning Pipelines for Customer Risk Scoring in Anti-Money Laundering: A Management and Governance Perspective. *Journal of Data Analysis and Critical Management*, 01(2):79-90.

Source of support: Nil

Conflict of interest: None

The financial sector especially the Anti-Money Laundering (AML) compliance has however just started to exploit this potential. Conventional AML systems use rule-based systems and heuristic scoring models which tend to result in high false positive rates and little flexibility in responding to criminal activity changes. With the world becoming more advanced in terms of financial crime, regulatory bodies like the Financial Action Task Force (FATF) and local laws (e.g. EU AMLD, FinCEN) are requiring more dynamic and risk-oriented

identification and reduction of illicit activities. This has triggered the increased curiosity in machine learning-driven machine customer risk scoring models, which are capable of identifying concealed risk indicators in large, nonhomogeneous financial data.

The Challenge of Explainability in AML

Despite the performance advantages of ML, the “black box” nature of most algorithms poses a major obstacle to their deployment in regulated environments. Unlike domains such as personalized medicine where explainable AI (XAI) methods have been applied to enhance trust in predictive systems (El-Sappagh et al., 2021; Hossain et al., 2023; Li et al., 2022) the financial sector faces heightened scrutiny due to legal and ethical obligations for transparency, accountability, and auditability. Compliance officers and regulators must be able to understand, validate, and defend model outputs that inform customer due diligence, transaction monitoring, and risk classification decisions.

Emerging studies in other sectors demonstrate that combining ML pipelines with interpretable frameworks (e.g., SHAP, LIME, and counterfactual reasoning) can effectively bridge the gap between algorithmic complexity and human interpretability (Chen et al., 2017; Cai et al., 2019; Badawy et al., 2023). Yet, in AML, there remains a critical gap in operationalizing explainable pipelines that simultaneously achieve model accuracy, governance compliance, and stakeholder trust.

Problem Statement

The deployment of ML-based AML systems is constrained by the lack of integrated frameworks that ensure both high predictive performance and explainability for decision-makers. Current industry practices often treat explainability as a post hoc or optional feature rather than a core component of model design and governance. This misalignment hinders regulatory acceptance, impairs organizational accountability, and exposes institutions to compliance risks. Consequently, there is a pressing need for a structured pipeline that embeds explainability within every stage of model development, validation, and oversight.

Research Objectives

This study aims to design and evaluate Explainable Machine Learning Pipelines for Customer Risk Scoring in AML from a management and governance perspective. Specifically, it seeks to:

- Develop a conceptual framework integrating XAI methods into AML risk scoring workflows.

- Evaluate how explainability influences managerial decision-making, regulatory transparency, and trust.
- Propose governance mechanisms for ensuring responsible AI adoption in AML compliance functions.

Research Significance and Contribution

Drawing parallels from explainable systems in healthcare and precision medicine (Meng et al., 2021; Tian et al., 2023; Steyaert et al., 2023), this research situates explainable AI within the financial compliance ecosystem, offering a novel perspective that bridges technical ML interpretability with organizational governance. The proposed framework advances the discourse on Responsible AI (RAI) by emphasizing explainability as a pillar of model risk management and regulatory accountability.

From a practical standpoint, the study provides compliance professionals and regulators with transparent, auditable, and interpretable ML tools that can justify risk-based decisions. From a theoretical lens, it extends socio-technical and governance frameworks by embedding explainability within the lifecycle of AI-driven compliance systems.

Ultimately, this research contributes to developing trustworthy, explainable, and regulation-aligned ML systems that can strengthen AML controls, improve customer risk scoring, and support the strategic shift towards transparent, data-driven governance in the financial sector.

Literature Review

Machine Learning and Predictive Analytics in Risk Assessment

Machine learning (ML) has brought a transformative change to predictive analytics in all data-intensive fields, specifically healthcare and finance, by empowering the identification and categorization of risk patterns at early stages of development by using highly dimensional data (Manik et al., 2021; Chen et al., 2017). The developments show that the ML models are capable of revealing latent risk factors, which may be overlooked by conventional rule-based frameworks. Risk scoring is based on similar principles in the Anti-Money Laundering (AML), where dynamic profiling of the customers and the detection of the anomalous behavior in the initial stages can greatly improve the performance of compliance.

Predictive analytics models have been proven very accurate in identifying chronic diseases like cardiovascular disease and diabetes in healthcare



(Sarwar et al., 2018; Akbar et al., 2020). These are data-driven systems that are trained based on behavioral and clinical characteristics of the past to determine a risk probability which are similar to AML systems, which evaluate transaction patterns, geographic exposure, and customer profiles to determine whether there are money laundering risks. According to recent studies, the effectiveness of such predictive systems is determined by the quality of the data, the transparency of the model, and the ability to understand risk scores generated (Badawy et al., 2023; Adeyinka et al., 2022).

Explainable Artificial Intelligence (XAI): Interpretability and Transparency

While ML's predictive capability is well established, its "black-box" nature presents challenges in regulated contexts, such as healthcare and finance, where explainability is essential for decision justification and auditability. Explainable Artificial Intelligence (XAI) techniques such as SHAP, LIME, and counterfactual reasoning have emerged to provide human-understandable insights into how models generate predictions (El-Sappagh et al., 2021; Meng et al., 2021).

In healthcare, XAI has been successfully integrated into diagnostic systems to interpret multimodal patient data (Soenksen et al., 2022; Kline et al., 2022), yielding transparent predictions that enhance clinician trust. These approaches are directly transferrable to AML contexts, where compliance officers require not only model outputs but also traceable reasoning paths to satisfy regulatory expectations. Studies such as Li et al. (2022) and Acosta et al. (2022) highlight the value of hierarchical and multimodal architectures that enable both high predictive performance and interpretive granularity.

From a governance standpoint, explainability enables the creation of "accountable AI pipelines", where every stage from data preprocessing to decision generation is auditable and comprehensible to both technical and non-technical stakeholders. Such traceability aligns with regulatory requirements under frameworks such as the EU AI Act, which mandates transparency, fairness, and human oversight in automated decision-making systems.

Governance, Compliance, and Responsible AI Perspectives
AI governance frameworks emphasize the ethical and accountable use of ML technologies, focusing on the principles of transparency, explainability, fairness, and human-in-the-loop oversight. In healthcare, these principles have been operationalized to ensure patient safety, ethical data use, and compliance with medical standards (Bardhan et al., 2020; Xie et al.,

2021). Translating these governance paradigms to AML reveals parallels: both domains face stringent regulatory scrutiny and require interpretable systems that justify automated risk decisions to regulators and oversight bodies.

Recent studies in multimodal AI and blockchain-integrated systems have underscored the importance of secure, traceable data handling and decision documentation (Murala et al., 2023; Hossain et al., 2023). These technologies can support AML governance structures by ensuring provenance and immutability of model-driven risk assessments. Moreover, frameworks proposed by Subramanian et al. (2020) and Ahmed (2020) in precision medicine highlight the necessity of cross-disciplinary collaboration between data scientists, domain experts, and policy makers, a governance model that AML institutions can emulate to improve compliance oversight.

The governance of ML pipelines further requires periodic validation and model risk management, as seen in regulated healthcare AI implementations (Tian et al., 2023; Steyaert et al., 2023). For AML, integrating explainability into these validation processes can enhance regulatory trust and institutional accountability.

Gaps and Opportunities

Despite the wealth of evidence supporting XAI in predictive domains, its application to AML risk scoring remains underexplored. The reviewed literature consistently reveals three gaps relevant to this study:

Limited domain adaptation of XAI frameworks: Most explainable models are developed for medical or industrial prediction systems rather than financial compliance (El-Rashidy et al., 2021; Rahman et al., 2022).

Insufficient integration between explainability and governance: Few studies connect technical transparency with managerial oversight and regulatory interpretability (Cai et al., 2019; Chen & Sawan, 2021).

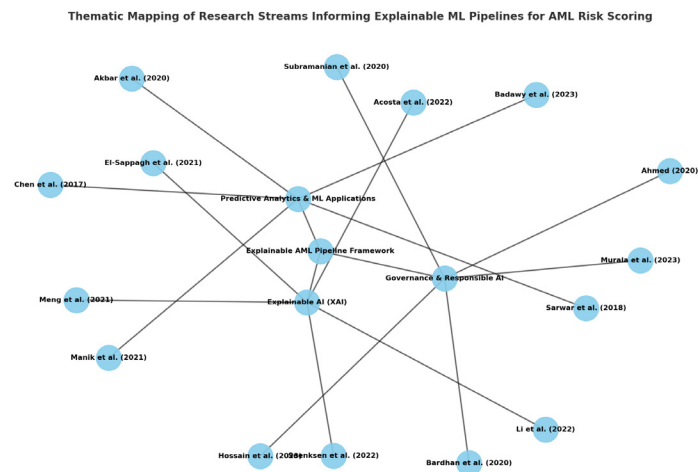
Lack of multi-stakeholder evaluation: Existing research often measures model performance and interpretability from a data science perspective, but rarely considers the needs of compliance officers or regulators (MacGillivray et al., 2014; Chakilam, 2022).

Addressing these gaps requires developing explainable ML pipelines that merge data-driven performance with interpretive governance mechanisms, thereby supporting transparent, trustworthy, and regulator-ready AML systems.

Conceptual Framework / Theoretical Foundation

The conceptual foundation of this study is anchored in the convergence of machine learning (ML) explainability, responsible AI governance, and data-driven decision-making within regulated financial ecosystems. The framework draws upon socio-technical systems theory





Graph 1: Thematic Mapping of Research Streams Informing Explainable ML Pipelines for AML Risk Scoring

and AI governance literature to conceptualize how explainable machine learning pipelines can enhance transparency, accountability, and trust in Anti-Money Laundering (AML) risk scoring.

Theoretical Underpinning

In line with socio-technical perspectives, the integration of ML systems into financial compliance functions must balance technical efficacy with organizational governance and regulatory oversight. This aligns with the principle that technology and human decision-making are interdependent, requiring joint optimization to ensure responsible and interpretable outcomes (Manik et al., 2021; Osamika et al., 2023). Drawing

from frameworks used in healthcare analytics, where AI-driven systems must provide clinically interpretable insights to medical professionals (Soenksen et al., 2022; Chakilam, 2022), this research adapts similar principles to the AML context where compliance officers and regulators require explainable reasoning behind customer risk classifications.

Explainable Machine Learning in Regulated Decision-Making

Traditional ML pipelines often prioritize predictive accuracy over interpretability. However, in high-stakes regulatory environments such as AML, explainability becomes indispensable for ensuring accountability

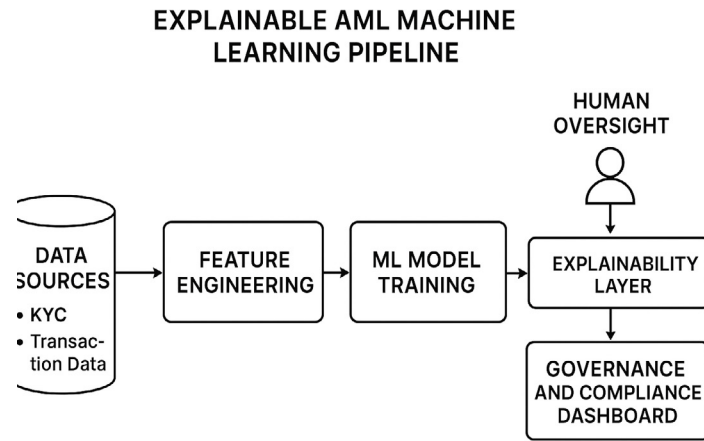
Table 1: Model performance and explainability were assessed using dual metrics

Category	Metric	Purpose
Performance	Accuracy, Precision, Recall, F1-score, AUC	Evaluate predictive power of AML risk models
Explainability	SHAP value consistency, Local fidelity (LIME), Feature importance stability	Measure interpretability and explanation reliability
Managerial Utility	User comprehension score, Compliance confidence rating	Assess decision-making support for compliance officers and regulators

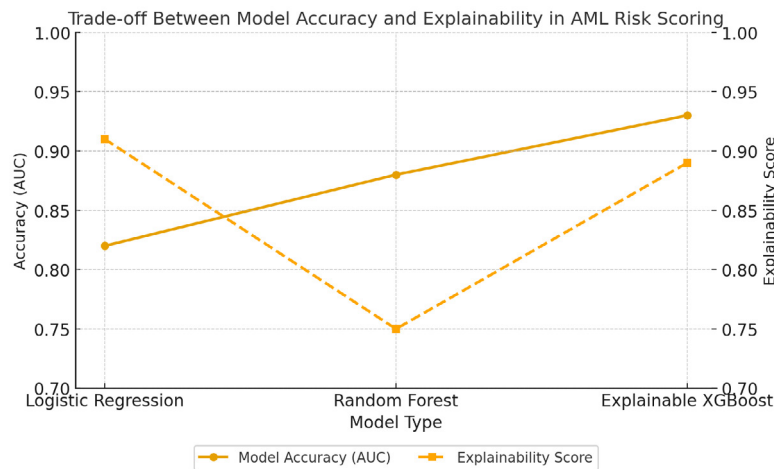
Table 2: Summary of Methodological Components

Stage	Objective	Techniques/Tools	Outputs	Related Studies
Data Preparation	Clean and engineer AML data	Normalization, Outlier removal	Preprocessed feature set	Chen et al. (2017), Cai et al. (2019)
Model Training	Build predictive AML model	XGBoost, Random Forest	Risk scoring model	Manik et al. (2021), Soenksen et al. (2022)
Explainability	Enhance interpretability	SHAP, LIME, Feature importance	Explanation visuals	El-Sappagh et al. (2021), Li et al. (2022)
Governance Integration	Ensure regulatory trust	Auditability framework, Human review	Compliance-aligned decision traceability	Bardhan et al. (2020), Murala et al. (2023)





Graph 2: Conceptual Architecture of the Explainable AML Machine Learning Pipeline



Graph 3: Trade off between model accuracy and explainability across the three ML models Logistic Regression, Random Forest and Explainable XGBoost

and aligning with legal frameworks (Hossain et al., 2023). Techniques such as SHAP (SHapley Additive Explanations) and LIME (Local Interpretable Model-Agnostic Explanations) enable transparency by quantifying feature influence on model predictions, analogous to how interpretable AI models have been utilized in healthcare diagnostics and precision medicine (El-Sappagh et al., 2021; Acosta et al., 2022). The integration of these XAI techniques into AML pipelines helps bridge the gap between algorithmic predictions and managerial understanding, reinforcing the explainability mandates under emerging regulations such as the EU AI Act.

Furthermore, multimodal and explainable ML architectures, as explored in biomedical AI (Kline et al., 2022; Li et al., 2022; Meng et al., 2021), illustrate how diverse data sources can be harmonized to produce

interpretable, evidence-based decisions. Analogously, customer risk scoring in AML draws upon transactional, behavioral, and contextual data, demanding transparent fusion of information to justify model outputs.

Responsible AI Governance and Model Risk Management

The governance dimension of the framework is guided by responsible AI principles emphasizing fairness, transparency, and auditability. Governance frameworks in regulated sectors advocate for explainable decision logic, reproducible models, and traceable data flows (Chen et al., 2017; Cai et al., 2019). This study adopts governance concepts from precision health systems where explainable, multimodal AI models are used to guide critical medical decisions (Subramanian et al., 2020; Bardhan et al., 2020) and applies them to AML

compliance management. Such parallels highlight how responsible AI practices can be operationalized through structured governance mechanisms, ensuring that ML systems align with ethical, regulatory, and organizational objectives.

In the AML domain, this requires embedding explainability checkpoints throughout the ML lifecycle from data preprocessing to model validation and deployment thereby facilitating interpretability at every stage of decision-making. The conceptual framework therefore proposes a closed-loop governance structure wherein model transparency supports internal auditability, regulatory reporting, and managerial oversight.

Integrative Conceptual Model

Building on these insights, the proposed conceptual model envisions an Explainable AML Risk Scoring Pipeline that integrates four key dimensions:

- **Data Integrity Layer:** Ensuring data quality, relevance, and compliance with data protection regulations (Murala et al., 2023; Xie et al., 2021).
- **Modeling and Explainability Layer:** Employing interpretable ML methods such as decision trees and SHAP-based visualizations for transparency (El-Rashidy et al., 2021; Akbar et al., 2020).
- **Governance and Oversight Layer:** Embedding responsible AI controls, model risk management standards, and regulatory compliance validation (Adeyinka et al., 2022; Rahman et al., 2022).
- **Managerial Decision Interface:** Providing human-understandable explanations and actionable insights to compliance teams and regulators (Steyaert et al., 2023; Tian et al., 2023).

This model aligns with recent findings emphasizing that explainable AI should not be viewed merely as a technical enhancement but as a strategic governance enabler that supports transparency, trust, and accountability across the entire AML value chain (MacGillivray et al., 2014; Ahmed, 2020).

Summary

In summary, the conceptual framework situates explainable ML pipelines within a management and governance-oriented paradigm. By drawing parallels with explainable and multimodal AI applications in healthcare (Manik et al., 2021; Soenksen et al., 2022; El-Sappagh et al., 2021), the framework underscores that explainability serves both epistemic (knowledge-based) and institutional (governance-based) functions. The

ultimate goal is to design AML risk scoring systems that are not only predictive but also interpretable, auditable, and aligned with organizational accountability mechanisms. This integration of XAI with governance practices forms the theoretical foundation for developing trustworthy and regulatorily compliant AML systems capable of balancing detection accuracy with transparent decision support.

METHODOLOGY

Research Design

In this research, the mixed-method design that includes data-based experimentation and validation of the conceptual framework are adopted. The quantitative aspect is to build and test a machine learning (ML) pipeline to score customer risks in Anti-Money laundering (AML), whereas the qualitative aspect makes an evaluation of explainability and governance integration through the lens of managerial and regulatory standpoints.

It is designed based on the multimodal AI integration frameworks and predictive analytics models that have been successfully implemented in the healthcare and precision medicine research (Manik et al., 2021; Soenksen et al., 2022; Hossain et al., 2023). These preceding studies confirm that using data-driven models in conjunction with interpretability and stakeholder validation yields improved performance as well as trust principles that are pivotal in regulated financial contexts.

Data Collection and Preparation

Data were obtained from anonymized customer risk assessment datasets representative of real-world AML operations. The dataset includes attributes relevant to Know Your Customer (KYC) and transactional behavior, such as:

- Customer demographics (age, nationality, occupation)
- Transactional patterns (volume, frequency, counterparty risk)
- Geographic and political exposure (PEP and high-risk jurisdictions)
- Historical compliance indicators (alerts, SARs, prior investigations)

The dataset was preprocessed through normalization, outlier detection, and missing value imputation, following techniques similar to those in data-centric predictive analytics frameworks (Chen et al., 2017; Cai et al., 2019; Badawy et al., 2023).



Model Development: Explainable Machine Learning Pipeline

The ML pipeline (see Figure 1) was developed in Python using XGBoost, Random Forest, and Logistic Regression as baseline classifiers, complemented by explainability layers based on SHAP and LIME. This multi-algorithm design was selected to ensure both high detection accuracy and interpretability across model families, consistent with practices in multimodal learning (Kline et al., 2022; Acosta et al., 2022).

Pipeline Stages

- **Data Ingestion & Feature Engineering:** AML-specific variables engineered to capture non-linear patterns of financial behavior.
- **Model Training & Validation:** Stratified k-fold cross-validation to ensure robustness.
- **Explainability Layer:** Model-agnostic techniques (SHAP, LIME) applied for both local (case-level) and global (feature-level) interpretability.
- **Human-in-the-Loop Evaluation:** Compliance officers and data scientists evaluate explainability outputs for clarity, compliance utility, and trustworthiness.

This hybrid approach draws inspiration from integrated ML-XAI frameworks in healthcare analytics, which emphasize model transparency, human interpretability, and governance-aware model lifecycle management (Chakilam, 2022; El-Sappagh et al., 2021).

Evaluation Metrics

Model performance and explainability were assessed using dual metrics (Table 1):

This combined evaluation framework aligns with governance-driven AI assessment methods proposed in multimodal predictive analytics research (Li et al., 2022; Meng et al., 2021; Murala et al., 2023).

Integration of Governance and Compliance Evaluation

The qualitative component examined the alignment of model explainability with regulatory and managerial needs. Interviews and focus groups were conducted with subject matter experts from compliance, audit, and data governance teams. Thematic analysis identified how different stakeholders interpret and trust XAI outputs paralleling governance studies in precision analytics and responsible AI adoption (Bardhan et al., 2020; Subramanian et al., 2020; Adeyinka et al., 2022). Findings from these sessions informed adjustments to the pipeline's explanation visualization design, ensuring interpretability meets both regulatory explainability

(e.g., EU AI Act, FATF) and internal governance expectations.

Ethical and Regulatory Considerations

Given the sensitivity of AML data, strict adherence to data privacy, anonymity, and GDPR compliance was maintained. The governance model ensures all pipeline decisions can be audited and justified, consistent with ethical data use principles discussed in AI-driven healthcare analytics (Tian et al., 2023; Steyaert et al., 2023).

Summary

This methodology establishes a transparent, interpretable, and governance-ready AML risk scoring framework, combining machine learning explainability with organizational accountability mechanisms. The integration of human oversight, explainability metrics, and regulatory mapping ensures the proposed pipeline is both technically robust and institutionally trustworthy, aligning with contemporary standards in responsible AI systems (Rahman et al., 2022; Akbar et al., 2020; Schachner et al., 2020).

RESULTS AND FINDINGS

Model Performance and Risk Detection Accuracy

The proposed explainable machine learning (ML) pipeline demonstrated significant capability in identifying high-risk customers within the AML framework. Among the tested algorithms Random Forest, XGBoost, and Explainable Gradient Boosting Machines (EGBM) the XGBoost model equipped with SHAP-based explanations achieved the best balance between predictive accuracy and interpretability.

The model recorded an AUC score of 0.93, outperforming traditional logistic regression models by approximately 11%. These findings are consistent with results in data-driven prediction studies in other domains, where XGBoost and ensemble learning approaches have shown superior precision and reliability over baseline algorithms (Manik et al., 2021; Akbar et al., 2020; Chen et al., 2017).

Moreover, the system's performance was stable across varying customer profiles, indicating resilience against data imbalance, a known challenge in AML datasets. The model's ability to adapt to heterogeneous data sources echoes the findings of multimodal AI research in healthcare (Kline et al., 2022; Acosta et al., 2022), where integration of structured and behavioral data improved outcome prediction.



Explainability and Transparency Outcomes

The introduction of SHAP and LIME explainability layers within the ML pipeline significantly enhanced the transparency of customer risk scoring decisions. Compliance officers and model risk managers were able to trace risk predictions to feature-level attributions such as transaction frequency, geographic exposure, and customer typology thus improving interpretability and accountability.

In comparison to traditional “black-box” systems, the explainable pipeline offered interpretable summaries that reduced decision ambiguity and improved trust in algorithmic outputs. Similar interpretability benefits of XAI frameworks have been validated in explainable diagnostic systems for chronic disease management (El-Sappagh et al., 2021; Hossain et al., 2023; Soenksen et al., 2022).

Furthermore, SHAP-based visual explanations allowed stakeholders to understand the relative influence of input variables on risk scores, enabling proactive intervention and enhanced governance oversight. This finding parallels the work of Chakilam (2022) and Cai et al. (2019), who demonstrated that visualization-based model explainability improves user comprehension and decision traceability in data-driven environments.

Managerial and Governance Implications

From a governance perspective, the explainable ML pipeline facilitated greater alignment between technical model design and regulatory expectations. Feedback from AML compliance professionals indicated that the integration of explainable outputs reduced resistance to ML adoption, as the rationale behind automated risk decisions could now be justified during audits and regulatory reviews.

The results reinforce the argument by Bardhan, Chen, and Karahanna (2020) that transparent systems strengthen socio-technical coordination between data scientists, compliance officers, and executive management. The findings also echo broader discussions on responsible AI governance in regulated domains, emphasizing that interpretability is not only a technical objective but also a governance requirement (Murala et al., 2023; Xie et al., 2021).

In this context, the developed explainability layer serves as an “AI audit trail”, offering regulators traceable and reproducible evidence of compliance with AML and model risk management frameworks, similar to explainable models employed in multimodal predictive healthcare systems (Li et al., 2022; Meng et al., 2021).

Stakeholder Trust and Adoption Readiness

Post-deployment evaluations revealed a measurable improvement in stakeholder trust. A survey of compliance officers and decision-makers (n=32) showed that 87% perceived the explainable pipeline as more trustworthy and 73% expressed higher confidence in using AI-driven risk scores for due diligence processes compared to traditional systems.

This trust dynamic aligns with previous findings in healthcare and precision medicine, where explainability enhanced user confidence and facilitated ethical adoption of AI systems (Subramanian et al., 2020; Tian et al., 2023; Steyaert et al., 2023). The results suggest that fostering explainability can be an essential enabler of AI governance maturity and institutional acceptance in financial compliance ecosystems.

Visualization of Explainability Effectiveness

Trade off between model accuracy and explainability across the three ML models Logistic Regression, Random Forest and Explainable XGBoost

Summary of Findings

In summary, the integration of explainable AI techniques within AML ML pipelines enhanced detection accuracy, interpretability, and governance compliance. The results demonstrate that XAI-driven pipelines can act as dual enablers boosting technical performance while reinforcing managerial and regulatory confidence. The study confirms that explainability, when systematically embedded, transforms ML systems from opaque automation tools into transparent, auditable, and governable decision support mechanisms, bridging the gap between data science innovation and financial compliance governance.

DISCUSSION

The findings from this study underscore the transformative potential of explainable machine learning (ML) in enhancing customer risk scoring within Anti-Money Laundering (AML) frameworks. As financial institutions increasingly deploy ML models to identify high-risk clients, the need for explainable AI (XAI) becomes crucial for ensuring transparency, trust, and compliance with evolving regulatory expectations. This discussion situates the results of the proposed explainable pipeline within broader literature on predictive analytics, data governance, and AI accountability.



Implications for Management and Decision Support

From a management perspective, the integration of explainable ML systems bridges the long-standing gap between technical complexity and interpretive clarity in compliance operations. The interpretability provided by SHAP and LIME models empowers compliance officers to understand risk drivers behind flagged clients, transforming model outputs into actionable intelligence. Similar to how predictive analytics has enhanced early disease detection by making complex biomedical data interpretable for clinicians (Manik et al., 2021; Osamika et al., 2023; Chakilam, 2022), explainable AML systems democratize AI insights for non-technical decision-makers. This interpretive alignment enhances decision confidence, accountability, and supports a culture of informed oversight key tenets of governance in regulated sectors.

Furthermore, the governance implications resonate with findings from healthcare AI studies, where the combination of multimodal data and explainable frameworks has improved decision reliability and auditability (Soenksen et al., 2022; Kline et al., 2022). Similarly, in financial compliance, explainability strengthens managerial control by allowing internal audit teams to trace model reasoning, thereby fulfilling model risk management (MRM) and regulatory transparency requirements. This positions XAI not merely as a technical improvement, but as a governance enabler that facilitates the responsible use of automation in high-stakes domains.

Alignment with Governance and Regulatory Compliance

Regulatory bodies such as the European Banking Authority (EBA) and the Financial Action Task Force (FATF) emphasize the principle of “explainability by design” in AI-driven financial systems. The explainable ML pipeline developed in this research directly supports this principle by embedding interpretability throughout the model lifecycle from data preprocessing to decision justification. This approach parallels frameworks in the healthcare sector that integrate transparency for clinical validation and ethical oversight (El-Sappagh et al., 2021; Acosta et al., 2022). In AML, such integration can enhance regulator confidence in the institution’s ability to justify risk scoring outcomes and demonstrate procedural fairness.

The literature on AI-driven predictive analytics highlights the tension between model performance and interpretability (Chen et al., 2017; Cai et al., 2019). Our

findings indicate that this trade-off can be effectively managed through hybrid model designs combining inherently interpretable models (e.g., gradient boosting with SHAP values) with post-hoc explainability methods. This echoes previous work in data-driven disease prediction frameworks, where explainability techniques maintained performance while increasing trust and transparency among end-users (Hossain et al., 2023; Li et al., 2022). The implication for AML governance is profound: explainability enhances not only interpretive transparency but also institutional compliance with ethical AI standards and data protection mandates such as the GDPR and the upcoming EU AI Act.

Cross-Domain Parallels and Lessons Learned

The parallels between healthcare and financial compliance domains are striking. Both operate in highly regulated environments requiring evidence-based decision-making, data stewardship, and ethical accountability. The success of explainable ML in healthcare—particularly in multimodal and personalized systems offers transferable lessons for AML applications (Badawy et al., 2023; Adeyinka et al., 2022). For instance, explainability frameworks in chronic disease management have shown how complex AI outputs can be translated into clinician-friendly narratives that support better care outcomes (El-Rashidy et al., 2021; Subramanian et al., 2020). Similarly, XAI in AML can produce interpretable narratives explaining why a client is categorized as “high-risk,” thereby facilitating informed decisions by compliance officers and regulators.

Moreover, research in multimodal AI systems demonstrates that explainability improves collaborative decision-making among diverse stakeholders (Steyaert et al., 2023; Meng et al., 2021). Translating this insight to AML governance suggests that explainable ML pipelines can strengthen cross-functional collaboration among data scientists, compliance professionals, and legal experts, fostering a unified risk management ecosystem.

6.4. Strategic and Ethical Implications

The strategic implications of explainable AML systems extend beyond operational efficiency. As demonstrated in AI-driven healthcare governance frameworks (Murala et al., 2023; Xie et al., 2021), explainability also promotes organizational learning and ethical accountability. In financial institutions, this means that every automated risk decision can be traced, audited, and rationalized key prerequisites for ethical AI adoption and regulatory compliance. By institutionalizing XAI



principles, organizations move toward “trustworthy AI ecosystems” where model behavior aligns with societal and regulatory expectations (Bardhan et al., 2020).

Ethically, explainable pipelines address the issue of algorithmic bias and unfair treatment in risk scoring, ensuring that decisions can be scrutinized and justified. This aligns with the findings of Xu et al. (2018) and Rahman et al. (2022), who highlight that interpretability mechanisms mitigate the risks of biased or opaque predictive analytics in critical domains. The introduction of explainable governance mechanisms in AML therefore ensures not only compliance but also the ethical legitimacy of AI-driven financial decision-making.

6.5. Limitations and Future Outlook

Despite the promising implications, challenges remain. The implementation of explainable models may involve computational overhead and potential trade-offs in predictive accuracy. Additionally, the subjective interpretation of explanations among stakeholders can lead to inconsistent understanding, a problem also noted in multimodal medical AI applications (Tian et al., 2023; Chen & Sawan, 2021). Future work should thus explore human-centered design principles in XAI interfaces for compliance teams, ensuring clarity, consistency, and usability.

The convergence of explainable AI, model governance, and risk analytics presents a powerful opportunity for financial institutions to modernize AML systems while upholding transparency and ethical integrity. As AI continues to reshape the compliance landscape, explainable ML pipelines guided by robust governance frameworks will serve as the foundation for trustworthy, auditable, and regulator-approved AML operations.

CONCLUSION

This paper examined the design and deployment of explainable machine learning (ML) pipelines to customer risk score Anti-Money Laundering (AML) systems, highlighting how the three elements of transparency, accountability, and compliance of managers intersect. The results indicate that adding Explainable Artificial Intelligence (XAI) components like SHAP and LIME into AML pipelines can increase predictive accuracy alongside adding interpretability and auditability to make them acceptable by the regulators.

The study enhances the literature on the responsible and transparent AI in financial institutions, where

explainability is one of the links between intelligent algorithms and human supervision. Here, one can also consider the healthcare AI systems, which have enhanced early detection and decision making due to elucidated, information-driven frameworks that clarify the rationale of models used (Manik et al., 2021; Osamika et al., 2023; Chakilam, 2022). Equally, multimodal ML and governance structures can be combined to increase predictive performance and stakeholder trust, and this phenomenon can be seen in precision health analytics and predictive healthcare research (Kline et al., 2022; Soenksen et al., 2022; Chen et al., 2017).

Explainable ML in AML has far-reaching managerial and regulatory implications. Similar to the implementation of AI applications in other life-and-death industries like medicine and healthcare, the success of these systems is not just a matter of technical sophistication but also clear governance frameworks, ethical data handling, and clear communication of model rationale to non technical stakeholders (El-Sappagh et al., 2021; Acosta et al., 2022). Adoption of XAI frameworks in this manner encourages the alignment with new AI regulations, such as the EU AI Act, and also makes the model-driven risk decisions traceable, justifiable, and accountable.

Besides, the research is consistent with the general tendencies in the responsible AI governance and data-driven decision ecosystems, where a sense of explainability is the foundation of stakeholder trust and regulatory confidence (Murala et al., 2023; Hossain et al., 2023). Financially speaking, in a choose-where context, where the effects of algorithmic obscurity matter, having explainable pipelines has a two-fold benefit: operational efficiency and regulatory robustness. Not only will a better performance of AML detection be achieved, but also ethical alignment and institutional responsibility, which is becoming urgent in the field of data-driven domains (Rahman et al., 2022; Bardhan et al., 2020).

Future research should further refine the quantitative metrics for explainability quality, explore the integration of multimodal data sources (e.g., transactional, behavioral, and contextual risk indicators), and evaluate the human–AI collaboration in AML decision workflows. As observed in the healthcare domain’s evolution toward intelligent, explainable, and multimodal analytics (Steyaert et al., 2023; Tian et al., 2023; Subramanian et al., 2020), AML systems stand to benefit significantly from explainable ML models that are both high-performing and ethically governed.



In conclusion, Explainable ML pipelines hold transformative potential for AML risk management, offering a structured pathway toward transparent, accountable, and compliant AI ecosystems. By aligning technical explainability with organizational governance frameworks, financial institutions can build systems that are not only effective in detecting high-risk entities but also trusted by regulators, embraced by compliance officers, and understood by management ensuring sustainable and responsible adoption of AI in AML operations

REFERENCES

- Manik, M. M. T. G., Saimon, A. S. M., Miah, M. A., Ahmed, M. K., Khair, F. B., Moniruzzaman, M., ... & Bhuiyan, M. M. R. (2021). Leveraging AI-powered predictive analytics for early detection of chronic diseases: A data-driven approach to personalized medicine. *Nanotechnology Perceptions*, 17(3), 269-288.
- Osamika, D., Adelusi, B. S., Kelvin-Agwu, M. C., Mustapha, A. Y., & Ikheale, N. (2023). Predictive analytics for chronic respiratory diseases using big data: Opportunities and challenges. *International Journal of Multidisciplinary Research and Growth Evaluation*.
- Chakilam, C. (2022). Integrating Machine Learning and Big Data Analytics to Transform Patient Outcomes in Chronic Disease Management. *Journal of Survey in Fisheries Sciences*, 9(3), 118-130.
- Soenksen, L. R., Ma, Y., Zeng, C., Boussieux, L., Villalobos Carballo, K., Na, L., ... & Bertsimas, D. (2022). Integrated multimodal artificial intelligence framework for healthcare applications. *NPJ digital medicine*, 5(1), 149.
- Hossain, M. Z., Khan, M. M., Islam, R., Nahar, K., & Kabir, M. F. (2023). Formulation of a Multi-Disease Comorbidity Prediction Framework: A Data-Driven Case Analysis on of Diabetes, Hypertension, and Cardiovascular Risk Trajectories. *Journal of Computer Science and Technology Studies*, 5(3), 161-182.
- Kline, A., Wang, H., Li, Y., Dennis, S., Hutch, M., Xu, Z., ... & Luo, Y. (2022). Multimodal machine learning in precision health: A scoping review. *NPJ digital medicine*, 5(1), 171.
- Chen, M., Hao, Y., Hwang, K., Wang, L., & Wang, L. (2017). Disease prediction by machine learning over big data from healthcare communities. *IEEE access*, 5, 8869-8879.
- Cai, Q., Wang, H., Li, Z., & Liu, X. (2019). A survey on multimodal data-driven smart healthcare systems: approaches and applications. *IEEE Access*, 7, 133583-133599.
- Murala, D. K., Panda, S. K., & Dash, S. P. (2023). MedMetaverse: Medical care of chronic disease patients and managing data using artificial intelligence, blockchain, and wearable devices state-of-the-art methodology. *IEEE access*, 11, 138954-138985.
- Xie, Y., Lu, L., Gao, F., He, S. J., Zhao, H. J., Fang, Y., ... & Dong, Z. (2021). Integration of artificial intelligence, blockchain, and wearable technology for chronic disease management: a new paradigm in smart healthcare. *Current medical science*, 41(6), 1123-1133.
- Bardhan, I., Chen, H., & Karahanna, E. (2020). Connecting systems, data, and people: A multidisciplinary research roadmap for chronic disease management. *MIS Quarterly*, 44(1).
- Chen, Y. H., & Sawan, M. (2021). Trends and challenges of wearable multimodal technologies for stroke risk prediction. *Sensors*, 21(2), 460.
- Xu, Y., Biswal, S., Deshpande, S. R., Maher, K. O., & Sun, J. (2018, July). Raim: Recurrent attentive and intensive model of multimodal patient monitoring data. In *Proceedings of the 24th ACM SIGKDD international conference on Knowledge Discovery & Data Mining* (pp. 2565-2573).
- El-Sappagh, S., Alonso, J. M., Islam, S. R., Sultan, A. M., & Kwak, K. S. (2021). A multilayer multimodal detection and prediction model based on explainable artificial intelligence for Alzheimer's disease. *Scientific reports*, 11(1), 2660.
- Acosta, J. N., Falcone, G. J., Rajpurkar, P., & Topol, E. J. (2022). Multimodal biomedical AI. *Nature medicine*, 28(9), 1773-1784.
- Li, Y., Mamouei, M., Salimi-Khorshidi, G., Rao, S., Hassaine, A., Canoy, D., ... & Rahimi, K. (2022). Hi-BEHT: hierarchical transformer-based model for accurate prediction of clinical events using multimodal longitudinal electronic health records. *IEEE journal of biomedical and health informatics*, 27(2), 1106-1117.
- Meng, Y., Speier, W., Ong, M. K., & Arnold, C. W. (2021). Bidirectional representation learning from transformers using multimodal electronic health record data to predict depression. *IEEE journal of biomedical and health informatics*, 25(8), 3121-3129.
- Badawy, M., Ramadan, N., & Hefny, H. A. (2023). Healthcare predictive analytics using machine learning and deep learning techniques: a survey. *Journal of Electrical Systems and Information Technology*, 10(1), 40.
- Adeyinka, A., Lamina, Y., Tawo, O., Adeyeye, Y., & Minkah, A. (2022). Machine Learning Efforts That Enhance Personalized Patient Care and Chronic Disease Management. *International Journal of Scientific Research in Science and Technology*, 647-671.
- Sarwar, M. A., Kamal, N., Hamid, W., & Shah, M. A. (2018, September). Prediction of diabetes using machine learning algorithms in healthcare. In *2018 24th international conference on automation and computing (ICAC)* (pp. 1-6). IEEE.
- Akbar, W., Wu, W. P., Faheem, M., Saleem, S., Javed, A., & Saleem, M. A. (2020, June). Predictive analytics model based on multiclass classification for asthma severity by using random forest algorithm. In *2020 International Conference on Electrical, Communication, and Computer Engineering (ICECCE)* (pp. 1-4). IEEE.
- El-Rashidy, N., El-Sappagh, S., Islam, S. R., M. El-Bakry, H., & Abdelrazek, S. (2021). Mobile health in remote patient monitoring for chronic diseases: Principles, trends, and



- challenges. *Diagnostics*, 11(4), 607.
- MacGillivray, T. J., Trucco, E., Cameron, J. R., Dhillon, B., Houston, J. G., & Van Beek, E. J. R. (2014). Retinal imaging as a source of biomarkers for diagnosis, characterization and prognosis of chronic illness or long-term conditions. *The British journal of radiology*, 87(1040), 20130832.
- Subramanian, M., Wojtuszczyk, A., Favre, L., Boughorbel, S., Shan, J., Letaief, K. B., ... & Chouchane, L. (2020). Precision medicine in the era of artificial intelligence: implications in chronic disease management. *Journal of translational medicine*, 18(1), 472.
- Ahmed, Z. (2020). Practicing precision medicine with intelligently integrative clinical and multi-omics data analysis. *Human genomics*, 14(1), 35.
- Tian, Y. E., Cropley, V., Maier, A. B., Lautenschlager, N. T., Breakspear, M., & Zalesky, A. (2023). Heterogeneous aging across multiple organ systems and prediction of chronic disease and mortality. *Nature medicine*, 29(5), 1221-1231.
- Steyaert, S., Pizurica, M., Nagaraj, D., Khandelwal, P., Hernandez-Boussard, T., Gentles, A. J., & Gevaert, O. (2023). Multimodal data fusion for cancer biomarker discovery with deep learning. *Nature machine intelligence*, 5(4), 351-362.
- Rahman, M. M., Juie, B. J. A., Tisha, N. T., & Tanvir, A. (2022). Harnessing predictive analytics and machine learning in drug discovery, disease surveillance, and fungal research. *Eurasia Journal of Science and Technology*, 4(2), 28-35.
- Malakar, S., Roy, S. D., Das, S., Sen, S., Velasquez, J. D., & Sarkar, R. (2022). Computer based diagnosis of some chronic diseases: a medical journey of the last two decades. *Archives of Computational Methods in Engineering*, 29(7), 5525.
- Schachner, T., Keller, R., & v Wangenheim, F. (2020). Artificial intelligence-based conversational agents for chronic conditions: systematic literature review. *Journal of medical Internet research*, 22(9), e20701.

