
DETECTING AND PREDICTING FINANCIAL FRAUD USING EXPLAINABLE AI & GRAPH-BASED ANOMALY DETECTION

¹Aditya Kudva and ²Tejashree Parab

Department of MSc Big Data Analytics, Jai Hind College (Autonomous), Mumbai, India

ABSTRACT

This literature paper review intends to highlight the shortcomings of traditional fraud detection methods, which often lack adaptability and transparency. Explainable AI (XAI) offers solutions by making model predictions more interpretable, enabling clearer insights for users. Additionally, graph-based anomaly detection captures complex relationships between financial entities, identifying patterns missed by traditional approaches. By integrating XAI with graph-based techniques, this review demonstrates improved accuracy and interpretability in fraud detection, offering advancements in financial auditing and compliance.

Keywords— *Financial statement fraud, Explainable AI (XAI), Graph-based anomaly detection, Fraud detection, Machine learning.*

INTRODUCTION

Financial statement fraud is a major concern for businesses and regulators, with traditional detection methods often lacking adaptability and transparency. While rule-based and black-box machine learning models are commonly used, they fail to evolve with changing fraud patterns and lack interpretability. Explainable AI (XAI) addresses this by making model predictions transparent, providing understandable insights for stakeholders. Additionally, graph-based anomaly detection captures complex relationships between financial entities, identifying irregular patterns indicative of fraud. This paper explores the integration of XAI and graph-based detection techniques to improve the accuracy and interpretability of financial fraud detection systems.

A. Machine Learning

Machine learning (ML) is key in detecting financial statement fraud by spotting complex patterns in data. In this research, ML is paired with Explainable AI (XAI) to make models more transparent and graph-based anomaly detection to catch sneaky fraud patterns. Traditional ML models, like decision trees or logistic regression, usually act like "black boxes" (mysterious, right?), but with XAI methods like SHAP and LIME, we get insights into why something looks fishy. Graph-based ML analyzes relationships between entities (companies, auditors, etc.), looking for weird patterns or unusual connections that scream fraud. This combo of ML, XAI, and graphs means better fraud detection and, more importantly, better explanations.

B. Explainable AI

Explainable AI (XAI) is crucial in this research for improving the transparency of fraud detection models. Traditional machine learning models often act as "black boxes," providing limited insights into their decision-making processes. XAI techniques, such as SHAP (SHapley Additive exPlanations) and LIME (Local Interpretable Model-agnostic Explanations), are used to make model predictions more understandable. By applying XAI, the model can explain why specific transactions or financial entities are flagged as fraudulent. This interpretability is essential for auditors, regulators, and stakeholders, as it increases trust in the model's output and ensures compliance with financial regulations. In combination with graph-based anomaly detection, XAI helps highlight suspicious patterns in financial data while making the detection process more transparent and actionable.

C. Anomaly Detection

In this research, anomaly detection serves as a foundational technique to identify instances of fraudulent behavior in financial data. It involves assessing the characteristics of various entities, such as companies and transactions, to pinpoint patterns that deviate from established norms. Traditional anomaly detection methods are employed, focusing on statistical approaches and machine learning algorithms to evaluate individual features of the data. The primary goal is to classify transactions or entities as either fraudulent or non-fraudulent based on their attributes. The implementation of anomaly detection in this context includes preprocessing steps such as data cleaning and feature engineering, which help enhance the quality and relevance of the data. By applying techniques such as the Iterative Imputer for handling missing values and SMOTE for addressing class imbalance, the research ensures that the dataset is prepared effectively for subsequent analysis. The results of the anomaly detection process inform the development of machine learning models, facilitating a comprehensive understanding of the features that contribute to fraud and improving predictive accuracy.

i) Graph Based Anomaly Detection

Graph-Based Anomaly Detection (GBAD) is a pivotal aspect of the research aimed at detecting financial fraud through the analysis of relationships among entities. By constructing a graph where nodes represent companies and transactions, and edges signify their interconnections, GBAD effectively captures the complex interactions that can indicate suspicious behavior. The research utilizes Euclidean distances to create a distance matrix, which is then transformed into a graph structure using NetworkX. The Asynchronous Label Propagation Algorithm is applied to identify communities within the graph, revealing clusters of behavior that may be indicative of fraud. Anomaly scores are calculated based on the sizes of these communities, with smaller communities suggesting higher potential for fraudulent activity. By setting a threshold at the 95th percentile of these scores, the research successfully flags entities for further investigation. This approach enhances traditional anomaly detection methods by uncovering hidden relationships and unusual patterns that are often missed in standard analyses. The integration of GBAD with machine learning models significantly contributes to the robustness and effectiveness of the fraud detection system developed in this research.

D. Predictive Analysis

In this research, predictive analytics is employed through a structured approach. The dataset undergoes preprocessing, including the handling of missing values with an Iterative Imputer and the creation of key financial ratios like Debt-to-Equity and Return on Equity. A binary target variable, Fraudulent, is established based on net income conditions. To address class imbalance, SMOTE is implemented, generating synthetic data to enhance model performance. Multiple classifiers, including XGBoost, Random Forest, and Logistic Regression, are trained and optimized via GridSearchCV, focusing on metrics such as F1-score. Explainability techniques like SHAP and LIME offer insights into feature importance and individual predictions, ensuring transparency in the models. Additionally, graph-based anomaly detection is utilized, scoring anomalies through community detection in a constructed graph of pairwise distances. This multifaceted predictive analysis aims to achieve accurate fraud detection while emphasizing interpretability, which is vital for stakeholder trust in automated financial systems.

E. Technical analysis

Technical analysis is pivotal for data understanding and model evaluation. The process begins with data preprocessing, which includes addressing missing values with the Iterative Imputer and creating financial ratios to enhance feature representation. You utilize various machine learning algorithms, notably XGBoost, Random Forest, and Logistic Regression, with hyperparameter tuning via GridSearchCV to optimize performance metrics such as F1-score. Class imbalance is managed through SMOTE, which generates synthetic samples, enhancing model robustness. To interpret model behavior, SHAP and LIME are employed to provide insights into feature contributions, promoting transparency in predictions. Additionally, a graph-based approach is applied, using community detection to uncover relationships in the data. By constructing a graph from pairwise distances and scoring anomalies, this technique identifies patterns indicative of fraudulent activity. Together, these technical methods form a comprehensive framework for effective fraud detection, emphasizing the need for both predictive accuracy and interpretability in machine learning.

The following are a few of the tools used in analyzing the stock prices

i) Regression Analysis

In the context of detecting financial statement fraud, regression analysis serves as a fundamental statistical tool for modeling relationships between financial variables and predicting fraudulent behavior. This method involves establishing a mathematical framework to analyze how changes in independent variables (financial ratios, indicators, etc.) affect a dependent variable, typically the likelihood of fraud.

The following are some of its variants that are used in the paper

a. Ordinary Least Squares (OLS)

Ordinary Least Squares (OLS) regression is utilized to detect anomalies in financial statements by modelling the relationship between financial metrics and the occurrence of fraud. By minimizing the sum of squared residuals, OLS estimates parameters in a linear regression model, enabling analysts to assess how various financial ratios correlate with fraud risk. Analysts can explore the impact of key ratios, such as return on equity (ROE) and debt-to-equity (D/E), on fraudulent reporting likelihood, where significant deviations from expected values may indicate irregularities necessitating further investigation. Additionally, OLS helps identify outliers; substantial divergences between actual and predicted financial metrics may signal potential manipulation or fraud.

b. Logistic Regression

Logistic regression enhances traditional regression by modeling binary outcomes, making it ideal for fraud detection. This technique estimates the probability of an event, such as fraud versus non-fraud, based on multiple predictor variables. In this research, logistic regression is employed to classify companies as fraudulent or non-fraudulent using financial indicators. The model yields probability scores that facilitate binary classifications, enabling targeted investigations of high-risk companies. Additionally, the coefficients from logistic regression reveal the influence of each financial ratio on fraud probability; positive coefficients indicate that higher ratios correlate with increased fraud likelihood, while negative coefficients suggest a lower probability.

ii) Time Series Analysis

Time series analysis is a vital statistical technique employed in this research to detect irregular patterns and trends in financial data over time. It allows for the observation of how financial metrics evolve, aiding in the identification of anomalies that may indicate fraudulent activities. By monitoring historical performance, analysts can discern long-term trends, such as revenue growth or declining expenses, and recognize sudden deviations that raise concerns about potential fraud. Additionally, time series analysis helps identify seasonal patterns inherent in financial data, distinguishing normal fluctuations from genuine anomalies. This analysis also facilitates the detection of outliers—data points that significantly deviate from expected values, which may signal manipulation. Moreover, time series models like ARIMA enable forecasting of future values based on historical data, allowing for the comparison of actual outcomes with predictions, thus highlighting discrepancies that may suggest fraud. Understanding autocorrelation further enhances fraud detection by revealing predictable relationships over time. Ultimately, time series analysis significantly contributes to financial statement fraud detection by providing insights into the temporal dynamics of financial data, enhancing the capacity to uncover fraudulent activities and ensure the integrity of financial reporting.

a. ARIMA (AutoRegressive Integrated Moving Average)

ARIMA (AutoRegressive Integrated Moving Average) is a widely used statistical model for analyzing time-dependent data, capturing key aspects of financial time series, including trends, seasonality, and autocorrelation. In the realm of financial statement fraud detection, ARIMA aids in modeling the historical performance of crucial financial metrics, enabling analysts to forecast future values. By applying an ARIMA model to financial data, researchers can identify deviations from predicted trends. Notable discrepancies between actual observations and forecasted values may indicate potential manipulation or fraud, thereby necessitating further investigation.

b. Seasonal decomposition

Seasonal decomposition techniques are employed to break down time series data into its underlying components: trend, seasonality, and residuals. This analysis is instrumental in identifying abnormal fluctuations in financial metrics that may signal fraudulent behavior. By decomposing financial data, analysts can effectively isolate seasonal patterns from irregularities, providing a clearer understanding of normal fluctuations versus unusual spikes or drops. Such insights enhance the detection of potential fraud, as unexpected changes can be indicative of manipulative practices.

Both regression analysis and time series analysis provide powerful tools for modeling relationships and detecting anomalies in financial data. OLS and logistic regression are crucial for assessing how various financial ratios correlate with fraud, while ARIMA and seasonal decomposition help analysts identify and interpret patterns over time, providing further context for potential irregularities in financial statements. By leveraging these techniques, analysts can gain deeper insights into the financial health of organizations and identify instances of financial statement fraud more effectively.

F. Fundamental analysis

Fundamental analysis is essential in researching financial statement fraud detection, as it evaluates a company's intrinsic value through key financial metrics and ratios, such as revenue, earnings, cash flow, and debt levels. This approach helps identify discrepancies or anomalies that may suggest fraudulent activities. By analyzing fundamental aspects, researchers can uncover inconsistencies in financial statements that deviate from expected norms. Significant variations in key ratios, like the debt-to-equity ratio or return on equity, compared to industry benchmarks may indicate potential red flags. Additionally, fundamental analysis establishes a baseline for expected performance, allowing analysts to detect irregularities over time, such as unexplained spikes in revenue or sudden drops in expenses. Integrating fundamental analysis with advanced techniques, such as regression and time series analysis, enhances the ability to detect financial fraud, ensuring a comprehensive evaluation of a company's financial integrity and operational practices.

The following are some tools that are used in Fundamental Analysis

i) Financial Ratios

Financial ratios are essential metrics that provide insights into a company's financial health and operational performance. Ratios such as the debt-to-equity ratio assess a company's leverage and risk, indicating how much debt is used to finance its assets relative to shareholder equity. The return on equity (ROE) measures profitability by assessing how effectively a company generates profits from its equity. These ratios help identify discrepancies that may suggest manipulation or fraudulent reporting.

a) Debt-to-Equity Ratio (D/E)

The Debt-to-Equity Ratio is a crucial measure of a company's financial leverage, calculated as the ratio of total liabilities to shareholders' equity. A high D/E ratio indicates that a company is significantly funded by debt relative to its equity, suggesting increased financial risk. In the context of fraud detection, unusually high or rapidly increasing D/E ratios may raise concerns about the company's financial health, potentially indicating manipulative practices such as inflating earnings or understating liabilities to present a more favorable image to investors.

$$D/E = \text{Total Liabilities} / \text{Shareholders Equity}$$

b) Return on Assets (ROA)

The Return on Assets ratio evaluates how efficiently a company utilizes its total assets to generate net income. It is calculated by dividing net income by total assets. A notably lower ROA compared to industry peers can be a red flag, signaling operational inefficiencies or potential earnings manipulation. This ratio serves as an indicator of whether a company is accurately reporting its profitability relative to its asset base, making it a vital tool for identifying discrepancies in financial statements.

$$ROA = \text{Net Income} / \text{Total Assets}$$

c) Current Ratio

The Current Ratio assesses a company's ability to meet its short-term obligations using its short-term assets, calculated as current assets divided by current liabilities. A ratio below 1 indicates potential liquidity issues, suggesting that the company may struggle to pay off its short-term debts. Furthermore, significant fluctuations in the current ratio could signal possible manipulation of financial statements, as companies might adjust their reported assets or liabilities to present a healthier liquidity position than what truly exists.

$$\text{Current Ratio} = \text{Current Assets} / \text{Current Liabilities}$$

d) Return on Equity (ROE)

The Return on Equity ratio measures a company's ability to generate profit from its shareholders' equity, calculated by dividing net income by shareholders' equity. A consistently increasing ROE is typically viewed positively, indicating effective management of equity. However, an abrupt spike in ROE may indicate aggressive accounting practices or earnings management. This raises concerns about the authenticity of reported profits and the potential for financial statement manipulation, prompting further scrutiny of the company's financial practices.

$$ROE = \text{Net Income} / \text{Shareholder's Equity}$$

These ratios play a fundamental role in the research, providing key insights into a company's financial health and helping analysts identify potential red flags associated with financial statement fraud. By evaluating these metrics, researchers can better assess the likelihood of manipulative practices and take necessary investigative actions.

METHODOLOGY

The following are the Machine Learning techniques used in our reviewed papers

XGBoost

XGBoost, an ensemble learning technique based on gradient boosting, is pivotal in fraud detection due to its ability to handle complex financial datasets with high dimensionality. In this research, XGBoost is applied to classify companies as fraudulent or non-fraudulent based on various financial indicators. Its gradient boosting framework improves prediction accuracy by iteratively combining weak learners and correcting errors from previous iterations. XGBoost's built-in regularization also helps prevent overfitting, making it well-suited for fraud detection, where patterns may be subtle and diverse.

Random Forest:

Random Forest, another ensemble method, is utilized in this research to enhance classification performance. By building multiple decision trees and averaging their results, Random Forest reduces variance and improves

prediction reliability. In the context of fraud detection, Random Forest captures complex relationships between financial variables and fraud risk, as the trees in the forest collectively identify important patterns and anomalies in the data. This method is particularly useful when dealing with noisy or unbalanced datasets, as it naturally handles such challenges through its bagging approach.

SHAP (SHapley Additive exPlanations)

SHAP is employed in this research to interpret the outputs of complex models like XGBoost, offering a clear, feature-level explanation of why a particular company is classified as fraudulent or non-fraudulent. By calculating the SHAP values, each feature's contribution to the model's prediction is quantified, making it possible to understand the importance of specific financial variables. This method ensures transparency by attributing the probability of fraud to individual financial ratios, providing a global view of which features consistently influence fraud detection across the dataset. SHAP helps stakeholders, such as auditors and regulators, trust the machine learning model by clearly demonstrating how predictions are made.

LIME (Local Interpretable Model-Agnostic Explanations)

LIME is used in this research to generate local explanations for individual predictions, particularly focusing on explaining specific fraud detection cases. It approximates the machine learning model locally around the point of interest, showing which financial variables contributed most to a company's classification as fraudulent or non-fraudulent. LIME makes complex models more interpretable by providing a simpler, understandable explanation for each individual decision. This tool is especially valuable when investigating high-risk companies, as it allows auditors to see why a specific prediction was made and identify the key factors driving the fraud classification.

NetworkX:

NetworkX is utilized in this research to represent financial metrics and companies as nodes in a graph. Edges between nodes represent relationships or similarities between financial metrics, such as the Euclidean distance between their values. By leveraging graph structures, the research models complex relationships between companies' financial data, which would be difficult to capture using traditional tabular methods. NetworkX facilitates the creation and analysis of these graphs, enabling the detection of irregular patterns that could indicate fraud. It allows for the visualization and understanding of interconnected financial metrics and their role in predicting fraudulent behavior.

Community Detection Algorithms:

Community detection algorithms, such as Asynchronous Label Propagation, are applied to the graph structures created using financial data. These algorithms group nodes into communities based on similarities or connections. In the context of fraud detection, the assumption is that normal companies will cluster into similar communities, while anomalous or fraudulent companies will either belong to unusually small or large communities or exhibit patterns that deviate from the norm. Identifying these anomalous patterns within communities allows for the detection of potential fraud, as companies with financial behaviors that differ significantly from their peers are flagged for further investigation.

SMOTE (Synthetic Minority Over-sampling Technique):

SMOTE is used in this research to address the inherent class imbalance in fraud detection datasets, where fraudulent cases typically represent a small portion of the overall data. Without addressing this imbalance, machine learning models may become biased toward the majority class (non-fraud cases). SMOTE generates synthetic examples of the minority class (fraud cases) by creating new samples based on the feature space of the existing minority instances. This helps ensure that the models trained on the resampled dataset can learn more effectively from the minority class, improving the model's ability to detect fraudulent activities.

Iterative Imputer:

The Iterative Imputer is applied to handle missing data in the dataset, which is common in real-world financial datasets. Missing values are filled in by iteratively modeling them based on the relationships between the other variables in the dataset. In this research, Iterative Imputer ensures that missing financial metrics do not undermine the model's performance. By imputing missing values with reasonable estimates, the technique preserves the integrity of the dataset and allows for more accurate analysis and model training.

Permutation Importance

Permutation Importance is used in this research to assess the importance of individual financial features in predicting fraud. By randomly shuffling specific features and observing how the model's performance changes, it quantifies the impact each feature has on the predictive accuracy. In the context of financial statement fraud detection, Permutation Importance helps identify which financial metrics (such as debt ratios or asset returns)

play a critical role in the classification of fraudulent vs. non-fraudulent cases. In this research, Permutation Importance aids in refining the machine learning models by pinpointing the most influential financial variables, thereby enhancing the interpretability of the model and improving its ability to target key indicators of fraud.

DATASET REVIEW

The literature papers reviewed had the following data sets

i) The Zacks Fundamentals Collection

The dataset comprises 84 columns and 1,534 rows, capturing financial data for various entities across specific time periods. Each row represents a unique company, while the columns encompass essential financial metrics, ratios, and data points pertinent to financial analysis.

Key identification columns include CoNo (a unique company identifier), CoName, CoType (company category), Exchange (stock exchange), and FiscalYearEnd (fiscal year-end month). Financial metrics such as Revenue (Rev), Net Income, Earnings Per Share (EPS), Dividend Yield (DivYield), and Book Value Per Share provide vital insights into company performance. Operational performance is assessed through metrics like Operating Income, EBIT, and EBITDA.

The dataset also includes critical ratios: profitability ratios (e.g., Return on Equity (ROE), Return on Assets (ROA)), leverage ratios (Debt to Equity, Current Ratio, Quick Ratio), and valuation ratios (Price-to-Earnings (PE), Price-to-Sales, Price-to-Book), which indicate market valuation relative to earnings, sales, and book value. Additionally, market performance data such as 52-Week High, 52-Week Low, Stock Price, and Volume offer a view of each company's market standing, while Beta measures volatility for stock analysis.

This rich dataset serves as a valuable resource for conducting fundamental analysis, enabling researchers and analysts to derive insights into company performance and market trends.

DATA PREPROCESSING

Prior to utilizing the dataset for analytical tasks, several preprocessing steps may be necessary to enhance data quality and model performance.

First, missing values can be an issue, as certain columns may contain incomplete data. To address this, imputation techniques such as the Iterative Imputer may be employed, or alternatively, rows or columns with substantial missing data can be removed to maintain dataset integrity.

Next, feature scaling is essential, given the wide range of financial metrics present in the dataset. Scaling is particularly important for machine learning models, which can be sensitive to the magnitude of input features.

Lastly, if the dataset is intended for fraud prediction analysis, it may exhibit class imbalance, characterized by a significantly lower number of fraudulent companies compared to non-fraudulent ones. To mitigate this issue, techniques like Synthetic Minority Over-sampling Technique (SMOTE) can be utilized to balance the class distribution, thereby improving model robustness and predictive accuracy.

These preprocessing steps are critical to ensure that the dataset is ready for effective analysis and modeling.

Implementation

The implementation for detecting and predicting financial fraud using Explainable AI and Graph-Based Anomaly Detection involves several systematic steps that integrate data preprocessing, exploratory data analysis, feature engineering, and advanced modeling techniques.

Initially, the dataset is loaded using pandas, followed by an exploration phase where the first few rows are displayed, and missing values are assessed. A heatmap visualization highlights the distribution of missing values, ensuring awareness of data quality before analysis.

To enhance the dataset's completeness, missing values are addressed using the Iterative Imputer, which employs a model-based approach to fill in gaps. Additionally, new financial ratios, such as Debt-to-Equity and Return on Equity, are engineered to provide deeper insights into company performance. A binary target variable, Fraudulent, is created based on a specified condition, such as negative net income, to facilitate fraud detection. Exploratory Data Analysis (EDA) follows, utilizing various visualizations like count plots and correlation matrices to understand feature distributions and relationships. A pairplot illustrates the characteristics of fraudulent and non-fraudulent entities, offering valuable insights into the dataset. Given the potential class imbalance—where fraudulent cases are significantly fewer than non-fraudulent ones—Synthetic Minority Over-sampling Technique (SMOTE) is employed. This technique generates synthetic samples to balance the dataset, ensuring robust model training. For the anomaly detection component, the code calculates

pairwise distances between samples to create a graph representation of the data. Community detection algorithms, such as asynchronous label propagation, are applied to identify anomalies. Anomaly scores are computed based on community sizes, and a threshold is established to classify potentially fraudulent companies. Model training is conducted using the XGBoost classifier, with hyperparameter tuning facilitated through GridSearchCV to optimize model parameters like the number of estimators and learning rate. To enhance model interpretability, SHAP (SHapley Additive exPlanations) values are utilized to explain the contributions of features to the model's predictions. Additionally, LIME (Local Interpretable Model-agnostic Explanations) provides local explanations for individual predictions, allowing for a detailed understanding of specific instances. Feature importance is further evaluated using permutation importance, which assesses the significance of each feature in predicting the target variable, and the results are visualized in a bar plot for clarity. The performance of the models is thoroughly evaluated using metrics such as accuracy, precision, recall, F1-score, ROC-AUC, and balanced accuracy for the XGBoost classifier, Random Forest, and Logistic Regression. A Voting Classifier is also implemented, combining the predictions from all models to leverage their strengths, ultimately enhancing overall performance. This comprehensive implementation effectively integrates multiple methodologies, combining traditional machine learning approaches with advanced graph-based techniques to enhance the detection of fraudulent financial activities. The incorporation of Explainable AI tools like SHAP and LIME ensures transparency in model predictions, which is crucial for stakeholders in the financial sector. Overall, the systematic approach to preprocessing, feature engineering, model training, and evaluation establishes a robust framework applicable to real-world financial datasets.

RESULTS

Table 1 Financial Metrics

Metric	Value
Debt_to_Equity	0.084718
Return_on_Equity	274.96786

Table 2 Model Evaluations

Model	Accuracy	Precision	Recall	F1 Score	Roc Auc
XG BOOST	1.000	1.0	1.0	1.0	1.00
Random forest	1.000	1.0	1.0	1.0	1.00
Logistic Regression	0.931	0.0	0.0	0.0	0.89
Voting Classifier	1.000	1.0	1.0	1.0	1.00

The company exhibits a solid financial foundation, highlighted by a low debt-to-equity ratio of 0.0847, indicating conservative financing primarily through equity, which minimizes financial risk. Additionally, the exceptionally high return on equity of 274.97% reflects the company's efficiency in generating profit relative to shareholders' equity, making it attractive to potential investors, though the sustainability of such returns warrants further examination. In terms of model evaluations for fraud detection, XGBoost, Random Forest, and Voting Classifier demonstrate perfect performance across all metrics, achieving an accuracy, precision, recall, F1 score, and ROC AUC of 1.000, effectively identifying fraud cases without errors. Conversely, Logistic Regression shows notable deficiencies, with an accuracy of 0.931 but a precision and recall of 0.0, indicating failure to predict any positive fraud cases. These insights underscore the company's financial stability while highlighting the efficacy of advanced machine learning models for effective fraud detection.

CONCLUSION

In this research, we have conducted a comprehensive analysis of financial statement fraud detection, employing various advanced machine learning models to effectively identify fraudulent activities within financial data. Our exploration began with a thorough examination of financial metrics, revealing a robust financial standing characterized by a low debt-to-equity ratio and an exceptionally high return on equity, indicating efficient profit generation with minimized financial risk.

We then focused on implementing and evaluating multiple machine learning models, including XGBoost, Random Forest, Logistic Regression, and a Voting Classifier. The results demonstrated that XGBoost, Random Forest, and the Voting Classifier achieved outstanding performance with perfect scores across all evaluation metrics, showcasing their effectiveness in accurately detecting fraud cases. In contrast, Logistic Regression

highlighted significant limitations, failing to identify any fraudulent instances despite a reasonable overall accuracy.

The challenges addressed in this research included the need for reliable fraud detection methods in financial statements, particularly in a landscape where traditional approaches may fall short. By leveraging explainable AI and graph-based anomaly detection techniques, we have developed a more nuanced understanding of fraud patterns, contributing to more effective prevention strategies. Overall, this study underscores the importance of integrating sophisticated machine learning techniques into financial analysis to enhance fraud detection capabilities, ultimately fostering greater transparency and trust in financial reporting. Future work may explore the sustainability of these findings and the applicability of the models across diverse datasets and industries.

REFERENCES

- [1] E. F. Brigham and M. C. Ehrhardt, *Financial Management: Theory & Practice*, 15th ed. Cengage Learning, 2016.
- [2] S. A. Ross, R. W. Westerfield, and B. D. Jordan, *Fundamentals of Corporate Finance*, 12th ed. McGraw-Hill Education, 2019.
- [3] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*. MIT Press, 2016.
- [4] C. M. Bishop, *Pattern Recognition and Machine Learning*. Springer, 2006.
- [5] D. D. M. B., "A survey of machine learning techniques for financial fraud detection," *International Journal of Data Mining & Knowledge Management Process (IJDMP)*, vol. 10, no. 2, pp. 1–15, 2020.
- [6] N. V. Chawla and A. Davis, "Improved data mining techniques for fraud detection," *Journal of the American Statistical Association*, vol. 97, no. 457, pp. 427–431, 2002.
- [7] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 2016, pp. 785–794.
- [8] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [9] D. R. Cox, "The regression analysis of binary sequences," *Journal of the Royal Statistical Society: Series B (Methodological)*, vol. 20, no. 2, pp. 215–232, 1958.
- [10] M. Sokolova and G. Lapalme, "A systematic analysis of performance measures for classification tasks," *Information Processing and Management*, vol. 45, no. 4, pp. 427–437, 2009.
- [11] D. M. W. Powers, "Evaluation: From precision, recall and F-measure to ROC, informedness, and correlation," in *Proceedings of the 23rd International Conference on Machine Learning (ICML)*, 2011, pp. 1–5.
- [12] M. T. Ribeiro, S. Singh, and C. Guestrin, "Why should I trust you? Explaining the predictions of any classifier," in *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 2016, pp. 1135–1144.
- [13] M. Ahmed, A. N. Mahmood, and J. Hu, "A survey of network anomaly detection techniques," *Journal of Network and Computer Applications*, vol. 60, pp. 120–137, 2016.
- [14] X. Wu and H. Liu, "Graph neural networks for anomaly detection in financial fraud," *ACM Transactions on Intelligent Systems and Technology (TIST)*, vol. 12, no. 5, pp. 1–25, 2021.
- [15] S. Bhattacharyya, S. Jha, and S. Tharakaram, "Data mining for fraud detection: A case study in banking," *International Journal of Computer Applications*, vol. 33, no. 7, pp. 1–5, 2011.
- [16] W. Wu and J. Hu, "Financial fraud detection: A survey from the data mining perspective," *Journal of Finance and Data Science*, vol. 1, no. 1, pp. 30–45, 2015.