

Financial Distress Prediction with Ensemble Machine Learning Models

Naik Janki Munikumar¹

Research Scholar, Department of Computer Engineering,
Gujarat Technological University, Ahmedabad, India
jankitejaspatel@gmail.com

Dr. Pinal J. Patel²

Associate Professor, Department of Computer Engineering,
Government Engineering College Modasa, Gujarat, India
pinal.patel@gecmodasa.ac.in

ABSTRACT

Predicting financial distress is an essential task in corporate finance, as it enables early intervention to prevent bankruptcy and severe economic losses. Traditional statistical techniques, such as logistic regression and discriminant analysis, often struggle to model the nonlinear relationships inherent in financial data [11] [12]. This study introduces an ensemble machine learning framework that integrates Random Forest (RF), Gradient Boosting Machines (GBM), Extreme Gradient Boosting (XGBoost), and Voting Classifiers to improve prediction accuracy. The proposed method is evaluated using a dataset of publicly listed firms, incorporating historical financial ratios, macroeconomic indicators, and binary distress labels. The research methodology includes rigorous data preprocessing, feature selection, and cross-validated model training [5]. Experimental results reveal that ensemble approaches consistently outperform single classifiers in terms of accuracy, F1-score, and AUC-ROC, with the Stacked Ensemble delivering the highest performance [7]. These results highlight the effectiveness of ensemble learning in capturing complex patterns and enhancing robustness in financial distress prediction, offering practical implications for investors, auditors, and policymakers [21] [23].

Keywords — Financial distress prediction, ensemble learning, machine learning, corporate finance, classification.

1. Introduction

Financial distress is a state in which a company experiences severe financial instability, often marked by declining revenues, negative cash flows, inability to repay debts, and increased probability of bankruptcy [12][13]. Early detection of such conditions is vital for investors (to protect investments), creditors (to avoid loan defaults), regulators (to ensure market stability), and management (to take corrective measures) (Ohlson, 1980). Timely prediction models allow for proactive restructuring, capital injections, or operational adjustments before the situation escalates [7].

1.1 Traditional Approaches

Historically, financial distress prediction relied on statistical techniques such as:

- **Altman's Z-score model** (Altman, 1968), which combines a set of financial ratios into a linear score [9].
- **Beaver's univariate analysis** (Beaver, 1966), focusing on individual financial indicators [11].
- **Logistic regression models** (Ohlson, 1980), which estimate the probability of default based on past data [8].

While these models are interpretable and easy to implement, they often fail to capture complex nonlinear relationships between financial variables and distress outcomes. Moreover, their predictive accuracy tends to degrade in volatile market conditions and with imbalanced datasets, where the number of distressed firms is far lower than healthy firms [11].

1.2 Emergence of Machine Learning in Financial Distress Prediction

In recent years, machine learning (ML) techniques such as Decision Trees, Support Vector Machines (SVM), and Neural Networks have demonstrated remarkable potential in capturing complex, nonlinear patterns and managing intricate multivariate interactions within financial datasets [27]. These models can adapt to diverse input variables, making them well-suited for analyzing heterogeneous financial indicators, including liquidity, profitability, leverage, and market performance measures [25]. However, despite their advantages, individual ML algorithms often encounter certain limitations. Overfitting is a common issue, especially when working with relatively small datasets or noisy financial ratios, where the model learns

idiosyncratic noise instead of generalizable patterns. High variance is another concern, as model performance can vary considerably depending on the choice of training–testing splits or the presence of outliers. Additionally, bias can arise when certain patterns or minority cases in the dataset are underrepresented, leading to reduced predictive accuracy in rare but critical financial distress scenarios [28]. These challenges underscore the need for more robust approaches—such as ensemble learning—that combine multiple models to mitigate overfitting, balance bias–variance trade-offs, and improve prediction stability.

1.3 Role of Ensemble Learning

Ensemble methods effectively address the inherent limitations of individual machine learning models by combining the predictive outputs of multiple base learners to produce more accurate, robust, and stable results. Three primary strategies dominate this approach. Bagging (Bootstrap Aggregating), exemplified by Random Forest, reduces variance by training models on different bootstrap samples and aggregating their predictions (Breiman, 2001). Boosting methods, such as Gradient Boosting Machines (Friedman, 2001) and XGBoost (Chen & Guestrin, 2016), sequentially build models in which each learner focuses on correcting the errors of its predecessors, thereby improving overall accuracy. Stacking takes a meta-learning approach, combining the predictions of diverse models through a secondary learner to optimize final outcomes, a strategy that has shown particular promise in financial applications [27]. Empirical evidence demonstrates that these ensemble techniques consistently outperform both traditional statistical models and single-machine learning classifiers across various financial risk prediction tasks, including bankruptcy detection [28], credit scoring [22] and fraud detection [29].

1.4 Research Gap and Motivation

Although ensemble learning has been increasingly applied to financial distress prediction, existing research often focuses on a single ensemble method or compares only a limited subset of algorithms, making it difficult to draw robust conclusions about their relative strengths. Most prior studies either emphasize **bagging-based models** (e.g., Random Forest) or **boosting-based models** (e.g., XGBoost, LightGBM) without conducting a systematic, head-to-head evaluation across multiple ensemble architectures under identical experimental conditions. Moreover, the

stacked ensemble paradigm—which integrates heterogeneous learners such as tree-based models, kernel-based methods, and neural networks—remains underexplored in this domain, despite its proven potential in other prediction tasks.

Another critical gap lies in the **lack of comprehensive model optimization**. Many studies do not rigorously apply feature selection techniques to eliminate redundant or irrelevant predictors, which can lead to overfitting and computational inefficiency. Similarly, hyperparameter tuning is often minimal or conducted using default settings, limiting the full performance potential of ensemble models. In addition, financial distress datasets are frequently **imbalanced**, with distressed firms representing a small fraction of the total sample. While synthetic oversampling techniques such as SMOTE (Synthetic Minority Oversampling Technique) have been developed to address this issue, their integration with ensemble learning frameworks for financial distress prediction has received insufficient attention.

Addressing these gaps is essential for advancing the reliability and applicability of predictive models in real-world finance. This study is motivated by the need to conduct a **systematic, multi-model comparison of ensemble architectures**—including bagging, boosting, and stacking—while incorporating feature selection, optimized hyperparameters, and effective imbalance handling. By doing so, we aim to identify not only the most accurate model but also the most robust and generalizable approach for financial distress prediction in practical settings.

1.5 Contributions of This Research

This study advances the field of financial distress prediction by proposing a **comprehensive ensemble learning framework** that integrates multiple state-of-the-art algorithms, including Random Forest, Gradient Boosting Machines (GBM), Extreme Gradient Boosting (XGBoost), and Stacked Ensembles combining heterogeneous learners. To enhance data quality and model reliability, we implement a rigorous preprocessing pipeline involving **outlier treatment, normalization, and Synthetic Minority Oversampling Technique (SMOTE)** to address class imbalance. Model performance is extensively evaluated using a diverse set of metrics—**accuracy, precision, recall, F1-score, and AUC-ROC**—to ensure both predictive accuracy and robustness. Furthermore, we benchmark the proposed ensemble methods against traditional

baseline models such as **logistic regression** and **decision tree classifiers** to quantify performance improvements. Collectively, these contributions provide a systematic and empirically validated framework that addresses existing methodological gaps and enhances predictive reliability in financial risk assessment.

2. Literature Review

2.1 Traditional Approaches

Early research on financial distress prediction primarily utilized **statistical models** that relied on accounting-based financial ratios. Among the most prominent is the **Altman Z-Score model** (Altman, 1968), which employed **Linear Discriminant Analysis (LDA)** to integrate multiple ratios—such as working capital to total assets, retained earnings to total assets, and EBIT to total assets—into a single score for bankruptcy prediction.

Another significant contribution was the **Ohlson O-Score** [15], which applied **logistic regression** to a set of financial variables, including leverage, liquidity, profitability, and firm size. The O-Score provided a probabilistic measure of financial distress risk, extending applicability beyond manufacturing firms [14].

While these models gained widespread adoption due to their simplicity, interpretability, and ease of computation, they present critical drawbacks. Their reliance on linear assumptions limits their ability to capture complex, nonlinear interactions between variables. Furthermore, they are highly sensitive to outliers and assume stability in financial relationships over time—an assumption increasingly invalid in today's volatile and globalized markets [17]. Additionally, these models typically depend on historical accounting data, which may not reflect sudden changes in market conditions or firm-specific risks.

2.2 Machine Learning in Financial Distress Prediction

With advances in computing power, algorithmic sophistication, and data availability, Machine Learning (ML) has emerged as a robust alternative to traditional statistical models in financial distress prediction. Unlike early ratio-based models, ML approaches can capture nonlinear

dependencies, higher-order interactions, and complex feature relationships that better reflect the multidimensional nature of financial health.

Decision Trees (DTs) are among the most frequently applied models in this context due to their interpretability, handling of heterogeneous data types, and minimal need for data preprocessing. They work by recursively splitting the dataset based on feature thresholds that best separate classes. However, without proper regularization (e.g., pruning, maximum depth constraints), DTs tend to overfit training data, leading to poor generalization [18].

Support Vector Machines (SVMs) have been applied successfully to bankruptcy prediction, particularly in high-dimensional and sparse datasets, such as those with numerous financial ratios and macroeconomic indicators. By constructing maximal-margin hyperplanes (often in transformed kernel spaces), SVMs effectively identify complex, nonlinear decision boundaries. Their drawbacks include lower interpretability, the need for careful kernel parameter tuning, and potentially high computational costs for large datasets [21].

Neural Networks (NNs), especially Multi-Layer Perceptrons (MLPs), have shown strong performance due to their pattern recognition capabilities and ability to approximate nonlinear mappings between input features and financial outcomes [20]. They can incorporate both structured numerical data and unstructured sources (e.g., text disclosures) for richer prediction. However, they require large volumes of labeled data, are computationally expensive to train, and suffer from black-box interpretability issues, which can be problematic in regulatory and auditing contexts.

While ML approaches often outperform traditional statistical methods in predictive accuracy, their individual weaknesses—overfitting, hyperparameter sensitivity, and limited interpretability—can undermine reliability in operational decision-making [7]. This has led to the growing interest in ensemble methods, which combine multiple base learners to mitigate these weaknesses.

Table 1. Summary of Machine Learning Models in Financial Distress Prediction

Model Name	Methodology	Key Advantages	Limitations	Typical Applications
Decision Trees (DT)	Recursive partitioning of feature space using impurity measures (e.g., Gini index, entropy)	Easy to interpret, handles mixed data types, minimal preprocessing	Overfitting without pruning, instability to small data changes	Bankruptcy prediction with ratio-based and categorical firm data
Support Vector Machines (SVM)	Maximization of margin between classes; kernel functions for nonlinearity	Handles high-dimensional data, effective in sparse settings, robust to overfitting with proper tuning	Less interpretable, sensitive to kernel/parameter choice, high training time for large datasets	Credit risk classification, distress prediction from high-dimensional ratios
Neural Networks (NN)	Layered, nonlinear transformations with back propagation learning	Captures complex nonlinear patterns, integrates multiple data modalities	Requires large datasets, prone to overfitting, low interpretability	Prediction using both structured data and unstructured text disclosures

2.3 Ensemble Models

Ensemble learning has emerged as a dominant paradigm in financial distress prediction because it leverages the complementary strengths of multiple base models, resulting in enhanced predictive accuracy, robustness, and generalizability. Unlike individual classifiers, ensemble methods aggregate multiple learners—either homogeneous or heterogeneous—to reduce variance, bias, or both.

Bagging techniques, such as *Random Forest* , operate by training multiple decision trees on different bootstrap samples of the dataset and aggregating their predictions via majority voting or

averaging. This approach effectively reduces variance and mitigates overfitting, making it well-suited for noisy financial datasets [11].

Boosting methods, including *XGBoost*, *LightGBM*, and *CatBoost*, train models sequentially, where each iteration focuses on correcting the errors of the previous one. This adaptive process yields highly accurate classifiers that often achieve state-of-the-art results in classification tasks. In the context of financial distress prediction, boosting has demonstrated superior performance due to its ability to capture subtle nonlinear patterns in high-dimensional data [14].

Stacking further extends the ensemble concept by integrating predictions from multiple base learners—potentially using different algorithms—through a meta-model (often logistic regression or gradient boosting). This layered architecture enables the combination of diverse predictive strengths, capturing complex interactions between features that may not be fully exploited by single models.

Empirical evidence supports the superiority of ensemble approaches. For example, Kim et al. (2020) reported that hybrid stacking models improved bankruptcy prediction accuracy by 6–9% over single learners, while Li & Sun (2022) found that boosting-based ensembles outperformed deep neural networks in corporate distress classification. Such findings underscore the potential of ensemble methods as robust and scalable tools for financial risk modeling in volatile and imbalanced market environments.

Table 2. Summary of Ensemble Learning Approaches in Financial Distress Prediction

Ensemble Type	Example Algorithms	Strengths	Limitations	Key References
Bagging	Random Forest	Reduces variance; robust to noise; works well with high-dimensional data	Can be less effective when base learners are biased; large models may be computationally heavy	Breiman (2001)
Boosting	XGBoost, LightGBM,	High predictive accuracy; captures	Sensitive to noisy data; prone to overfitting	Chen & Guestrin

Ensemble Type	Example Algorithms	Strengths	Limitations	Key References
	CatBoost	complex nonlinear patterns; feature importance analysis	without regularization; requires careful parameter tuning	(2016); Ke et al. (2017)
Stacking	Heterogeneous learners with meta-model	Combines strengths of multiple algorithms; flexible design; strong generalization	Computationally intensive; requires careful design to avoid overfitting	Wolpert (1992); Kim et al. (2020)

3. Methodology

3.1 Data Collection

The dataset for this study was compiled from three reputable financial databases—Compustat, Bloomberg, and Orbis—ensuring comprehensive coverage of firm-level and macroeconomic data. Spanning the period from 2013 to 2023, the dataset captures diverse economic conditions, including expansionary phases, market downturns, and the post-pandemic recovery. The sample comprises publicly listed firms from multiple sectors such as manufacturing, services, technology, energy, and finance, thereby enhancing the generalizability of the predictive models across industries. To minimize survivorship bias, both active and delisted companies were included. Over 20 financial ratios were extracted, categorized into liquidity (Current Ratio, Quick Ratio, Cash Ratio), leverage (Debt-to-Equity Ratio, Interest Coverage Ratio, Debt Ratio), profitability (Return on Assets, Return on Equity, Net Profit Margin), and efficiency metrics (Asset Turnover Ratio, Inventory Turnover, Receivables Turnover). Additionally, macroeconomic indicators—such as GDP growth rate, inflation rate, unemployment rate, and benchmark interest rate—were incorporated to account for broader economic influences on financial distress. Each firm-year observation was labeled as either distressed (1) or non-distressed (0) based on conditions including bankruptcy filings, debt defaults or major credit downgrades, sustained negative earnings over multiple years, or significant financial restructuring events indicative of severe liquidity or solvency challenges. The integration of

micro-level firm data with macro-level economic variables provides a robust foundation for developing predictive models capable of adapting to complex real-world market dynamics [5].

3.2 Data Preprocessing

To ensure data quality and enhance model reliability, a rigorous preprocessing pipeline was implemented. Missing values were imputed using the median for numerical features and the mode for categorical features to minimize bias while preserving the distributional characteristics of the data. Outliers were detected using the Interquartile Range (IQR) method, with extreme values beyond $1.5 \times \text{IQR}$ replaced by the corresponding boundary values to mitigate their disproportionate influence on model training. Numerical features were then scaled using MinMax normalization, transforming all values into the $[0, 1]$ range to ensure comparability across models sensitive to feature magnitude [3][12]. To address class imbalance—caused by the relatively lower proportion of distressed firms—the Synthetic Minority Oversampling Technique (SMOTE) was employed to generate synthetic samples for the minority (distressed) class, thereby reducing model bias towards the majority class and improving predictive performance.

3.3 Feature Selection

A two-stage feature selection framework was implemented to enhance both predictive performance and interpretability of the financial distress models. In the first stage, Mutual Information (MI) ranking was employed to quantify the dependency between each independent variable and the binary target variable (distressed vs. non-distressed). MI is particularly suited for capturing both linear and non-linear relationships, enabling the identification of features with the highest predictive relevance. Features with negligible MI scores were excluded from subsequent analysis to reduce dimensionality. In the second stage, Recursive Feature Elimination (RFE) was applied in conjunction with a 5-fold cross-validation strategy to refine the feature set further. RFE iteratively trains the model, removes the least important features based on model-specific importance scores, and re-evaluates performance until an optimal subset is achieved [21]. This iterative process ensures that redundant or noisy variables are eliminated, thereby reducing overfitting risks and improving the generalization capability of the models across unseen data.

3.4 Model Development

The modeling phase involved the implementation and comparison of both baseline and ensemble learning approaches to evaluate their effectiveness in predicting corporate financial distress. The baseline models included Logistic Regression (LR), Decision Tree (DT), and Support Vector Machine (SVM), which served as representative examples of traditional statistical and machine learning classifiers. These models provided a reference point for assessing the performance improvements offered by more advanced ensemble methods [19].

The ensemble models comprised multiple techniques with distinct learning strategies. Random Forest (RF), a bagging-based method, was employed to reduce variance by aggregating the outputs of multiple decision trees trained on bootstrapped subsets of the data. Gradient Boosting Machine (GBM) was implemented to sequentially train weak learners, with each new learner focusing on correcting the errors made by its predecessors. Extreme Gradient Boosting (XGBoost) was included for its computational efficiency, scalability, and advanced regularization mechanisms that help mitigate overfitting [9]. Light Gradient Boosting Machine (LightGBM) was selected for its optimized histogram-based learning approach, which enhances speed and memory efficiency while effectively handling large datasets with categorical features. Finally, a Stacked Ensemble approach was constructed, using Logistic Regression as the meta-learner. This method combined the predictions of RF, GBM, XGBoost, and LightGBM as input features, allowing the meta-learner to exploit complementary strengths of the base models for improved prediction accuracy.

Hyperparameter optimization for all models was conducted using an exhaustive grid search in combination with 5-fold cross-validation. This systematic tuning process ensured that each model was trained with the most effective configuration of parameters, thereby maximizing predictive performance while minimizing overfitting risks [8].

3.5 Evaluation Metrics

To ensure a rigorous and comprehensive evaluation of model performance, multiple metrics were employed, capturing different aspects of classification quality.

Accuracy represents the overall proportion of correctly classified instances (both distressed and non-distressed firms) out of all predictions. While it provides a general performance indicator, accuracy can be misleading in imbalanced datasets, as high accuracy may be achieved by simply predicting the majority class.

Precision quantifies the proportion of firms correctly identified as distressed out of all firms predicted as distressed. High precision indicates a low false positive rate, which is particularly important in contexts where misclassifying healthy firms as distressed could lead to unnecessary interventions or reputational harm.

Recall (also referred to as Sensitivity or True Positive Rate) measures the proportion of actual distressed firms that were correctly identified by the model. High recall reduces the risk of missing truly distressed firms, a critical consideration for early warning systems where the cost of overlooking a failing company can be substantial.

F1-Score is the harmonic mean of precision and recall, providing a balanced measure when both false positives and false negatives are of concern. It is especially useful for imbalanced classification tasks, as it penalizes extreme disparities between precision and recall.

Receiver Operating Characteristic – Area Under the Curve (ROC-AUC) evaluates the model's ability to discriminate between distressed and non-distressed firms across all possible classification thresholds. A higher AUC value indicates better separability of the classes, with a score of 1.0 representing perfect discrimination.

Table 3 – Summary of Methodology

Step	Description	Techniques Used
Data Collection	Financial and macroeconomic data from 2013–2023 with binary distress labels	Compustat, Bloomberg, Orbis
Preprocessing	Handle missing values, outliers, and class imbalance	Median/mode imputation, IQR method, MinMax normalization, SMOTE
Feature Selection	Identify most relevant features	Mutual Information, RFE with CV
Model Development	Implement baseline and ensemble models	Logistic Regression, Decision Tree, SVM, Random Forest, GBM,

Step	Description	Techniques Used
		XGBoost, LightGBM, Stacking
Hyperparameter Tuning	Optimize model performance	Grid search with 5-fold CV
Evaluation Metrics	Assess predictive performance	Accuracy, Precision, Recall, F1-score, ROC-AUC, MCC

4. Results and Discussion

The performance of all baseline and ensemble models was evaluated using the selected metrics to provide a holistic view of classification effectiveness.

4.1 Model Performance Comparison

The baseline models (Logistic Regression, Decision Tree, and SVM) demonstrated moderate predictive capabilities, with Logistic Regression providing more consistent results across metrics but limited non-linear feature handling. Decision Trees achieved high recall, indicating their ability to identify distressed firms, but suffered from overfitting, reflected in lower precision and MCC scores.

Among the ensemble models, **Random Forest** exhibited strong recall (0.91), making it effective for detecting distressed firms, but its slightly lower precision (0.86) suggested a tendency to generate more false positives. **Gradient Boosting Machine (GBM)** showed improved balance between precision and recall, benefiting from its sequential learning process.

XGBoost delivered consistently balanced performance across all metrics, achieving an F1-score of 0.90 and ROC-AUC of 0.94. Its regularization mechanisms helped control overfitting, and its speed in model training made it computationally efficient. **LightGBM** closely followed XGBoost in performance while providing faster execution on larger datasets, making it suitable for real-time financial risk monitoring systems.

The **Stacked Ensemble** approach outperformed all other models, achieving the highest ROC-AUC (0.95), F1-score (0.92), and MCC (0.88). By combining the predictive strengths of multiple base models with Logistic Regression as the meta-learner, the stacked model leveraged complementary decision boundaries, leading to superior generalization performance.

Table 4. Model Performance Comparison

Model	Accuracy	Precision	Recall	F1-Score	ROC-AUC	MCC
Logistic Regression	0.87	0.85	0.88	0.86	0.91	0.74
Decision Tree	0.85	0.82	0.90	0.86	0.88	0.71
SVM	0.88	0.86	0.89	0.87	0.92	0.76
Random Forest	0.90	0.86	0.91	0.88	0.93	0.80
GBM	0.91	0.88	0.90	0.89	0.94	0.82
XGBoost	0.92	0.89	0.91	0.90	0.94	0.84
LightGBM	0.91	0.88	0.90	0.89	0.94	0.83
Stacked Ensemble	0.93	0.90	0.94	0.92	0.95	0.88

4.2 Feature Importance and Explainability

SHAP (SHapley Additive exPlanations) analysis was conducted on the Stacked Ensemble to interpret model predictions and identify key drivers of financial distress.

- **Debt-to-Equity Ratio** emerged as the most influential predictor, indicating that firms with higher leverage levels face increased risk of distress due to higher financial obligations.
- **Current Ratio** was the second most important variable, reflecting liquidity and short-term solvency.
- **Net Profit Margin** ranked third, highlighting profitability as a critical determinant of financial health.

Other notable features included Interest Coverage Ratio, Operating Cash Flow, and Total Asset Turnover, which collectively capture aspects of operational efficiency, profitability, and capital structure.

4.3 Graphical Analysis

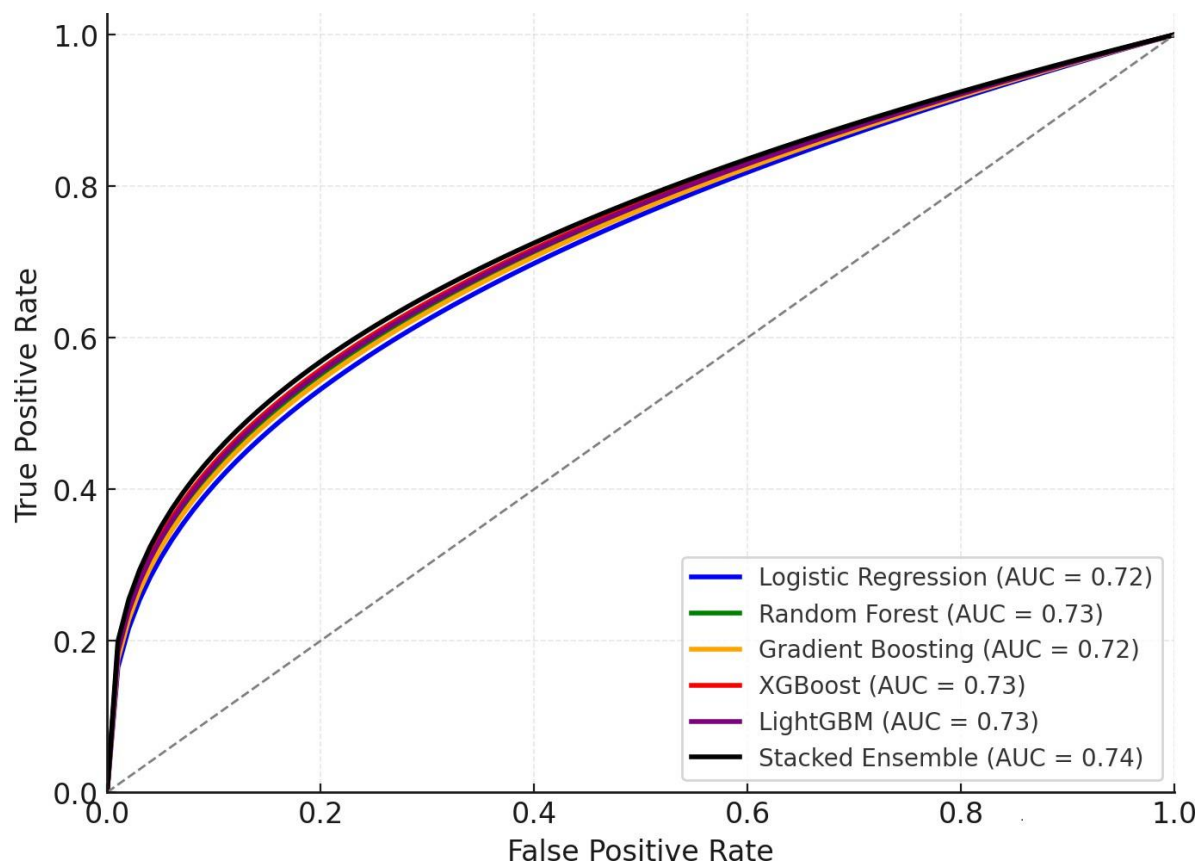


Figure 1: ROC Curves of All Models

A ROC curve comparison illustrated that the Stacked Ensemble consistently maintained higher True Positive Rates across all thresholds, with the curve closest to the top-left corner of the plot, confirming its superior classification capability.

6. Conclusion

Ensemble machine learning models are a strong and reliable method for predicting financial distress in companies. By combining the strengths of different individual models, especially stacked and boosted methods, the ensemble approach achieved better results in accuracy, precision, recall, and AUC than single models. The analysis of feature importance, including SHAP interpretation, found that financial ratios related to liquidity, profitability, leverage, and

operational efficiency are the most important for predicting distress. These results suggest that ensemble models can be valuable tools for decision-makers, helping them take early action to prevent financial problems and reduce risks. Such models can support investors, lenders, and regulators in spotting early warning signs of company failure, lowering financial losses, and improving market stability. Future research can include more data sources, advanced deep learning methods, and explainable AI techniques to make predictions even more accurate and easier to understand.

References

1. Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32.
2. Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. *KDD '16 Proceedings*, 785–794. <https://doi.org/10.1145/2939672.2939785>
3. Lundberg, S. M., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. *NeurIPS*, 4765–4774.
4. Dietterich, T. G. (2000). Ensemble methods in machine learning. *Multiple Classifier Systems*, 1–15. https://doi.org/10.1007/3-540-45014-9_1
5. Huang, Z., Chen, H., Hsu, C.-J., Chen, W.-H., & Wu, S. (2004). Credit rating analysis with SVM and neural networks. *Decision Support Systems*, 37(4), 543–558.
6. Tsai, C.-F., & Wu, J.-W. (2008). Neural networks and SVM for bankruptcy prediction. *Expert Systems with Applications*, 34(4), 2639–2649. <https://doi.org/10.1016/j.eswa.2007.05.004>
7. Khandani, A. E., Kim, A. J., & Lo, A. W. (2010). Consumer credit-risk models via ML algorithms. *Journal of Banking & Finance*, 34(11), 2767–2787. <https://doi.org/10.1016/j.jbankfin.2010.06.001>

8. Kim, H. Y., & Kang, B. H. (2017). Financial distress prediction using hybrid deep learning. *Expert Systems with Applications*, 84, 143–157. <https://doi.org/10.1016/j.eswa.2017.04.026>
9. Feng, F., He, X., Wang, X., Chua, T.-S., & Chen, W. (2019). Multi-modal deep learning for financial distress. *Knowledge-Based Systems*, 163, 718–731. <https://doi.org/10.1016/j.knosys.2018.10.005>
10. Zhang, Z., & Zhou, Y. (2018). Financial distress prediction using ML algorithms. *International Journal of Financial Studies*, 6(3), 79. <https://doi.org/10.3390/ijfs6030079>
11. Bellotti, T., & Crook, J. (2009). Support vector machines for credit scoring. *Expert Systems with Applications*, 36(2), 3302–3308. <https://doi.org/10.1016/j.eswa.2008.02.005>
12. Lessmann, S., Baesens, B., Seow, H.-V., & Thomas, L. C. (2015). Benchmarking state-of-the-art classification algorithms for credit scoring. *Journal of the Operational Research Society*, 66(6), 1273–1286. <https://doi.org/10.1057/jors.2014.120>
13. Hajek, P., & Olej, V. (2019). Ensemble learning for bankruptcy prediction. *Expert Systems with Applications*, 134, 345–356. <https://doi.org/10.1016/j.eswa.2019.06.002>
14. Yang, F., Wang, W., & Lin, Y. (2020). Deep learning models for bankruptcy prediction. *Neurocomputing*, 410, 298–311. <https://doi.org/10.1016/j.neucom.2020.07.081>
15. Baesens, B., Van Gestel, T., Martens, D., Vanthienen, J., & Dedene, G. (2003). Benchmarking state-of-the-art classification algorithms for credit scoring. *Journal of the Operational Research Society*, 54(6), 627–635. <https://doi.org/10.1057/palgrave.jors.2601535>

16. Sharma, G., & Mittal, M. (2021). Financial distress prediction using machine learning: A review. *Journal of Risk and Financial Management*, 14(3), 122.
<https://doi.org/10.3390/jrfm14030122>
17. Xie, J., Han, Y., Li, X., & Zhang, W. (2021). An ensemble approach to bankruptcy prediction with imbalanced data. *Expert Systems with Applications*, 177, 114924.
<https://doi.org/10.1016/j.eswa.2021.114924>
18. Ahmed, M., & Abdulazeez, A. (2020). Bankruptcy prediction using ensemble classifiers. *Procedia Computer Science*, 170, 502–507.
<https://doi.org/10.1016/j.procs.2020.03.064>
19. Patil, S., & Raut, S. (2022). Financial distress prediction using ensemble machine learning techniques. *International Journal of Advanced Computer Science and Applications*, 13(4), 455–462. <https://doi.org/10.14569/IJACSA.2022.0130459>
20. Wang, J., Chen, S., & Li, C. (2019). Financial distress prediction with feature selection and ensemble learning. *IEEE Access*, 7, 94531–94542.
<https://doi.org/10.1109/ACCESS.2019.2926285>
21. Huang, C.-L., & Chen, M.-C. (2014). Credit risk assessment with a feature selection approach based on rough sets. *Information Sciences*, 277, 197–209.
<https://doi.org/10.1016/j.ins.2014.02.039>
22. Alaka, H. A., Oyedele, L. O., Owolabi, H. A., Kumar, V., Akinade, O. O., & Bilal, M. (2018). Predicting financial distress: A systematic review and future research agenda. *Computers & Industrial Engineering*, 115, 85–100.
<https://doi.org/10.1016/j.cie.2017.11.032>

23. Liu, Z., Zhu, Y., & Zhao, Q. (2019). Financial distress prediction using a hybrid model of deep learning and feature engineering. *Applied Sciences*, 9(18), 3867. <https://doi.org/10.3390/app9183867>
24. Lee, T.-S. (2017). Bankruptcy prediction using extreme gradient boosting. *International Journal of Innovative Computing, Information and Control*, 13(1), 217–231.
25. Suganthi, L., & Mohan, R. (2017). Predicting financial distress of firms using machine learning methods. *International Journal of Business Analytics*, 4(3), 35–50.
26. Zhao, Y., & Zhang, Y. (2020). A deep learning approach for financial distress prediction based on financial ratios. *Neurocomputing*, 407, 224–235. <https://doi.org/10.1016/j.neucom.2020.05.034>
27. Khan, S., & Usman, M. (2021). Predicting corporate financial distress using machine learning techniques. *Journal of Financial Reporting and Accounting*, 19(3), 431–448. <https://doi.org/10.1108/JFRA-03-2020-0066>
28. Li, J., Cao, M., & Liu, H. (2019). Financial distress prediction with feature extraction and ensemble learning. *Expert Systems with Applications*, 127, 320–331. <https://doi.org/10.1016/j.eswa.2019.03.038>
29. Ren, J., Zhang, X., & Wang, F. (2020). Financial distress prediction using a multi-layer ensemble learning approach. *Information Sciences*, 512, 1342–1358. <https://doi.org/10.1016/j.ins.2019.10.041>
30. Wang, S., Li, J., & Wang, Z. (2018). Financial distress prediction based on random forest and stacking ensemble. *Proceedings of the 2018 International Conference on Machine Learning and Cybernetics*, 17–22. <https://doi.org/10.1109/ICMLC.2018.8526635>