

IRisk Assessment Model in Enterprise Investment and Financing Decision-making Based on Time Series Analysis

Fahad Alam*

School of Economics and Management, Wuhan University, Wuhan, China; Fahadalam@whu.edu.cn

*Corresponding Author: Fahadalam@whu.edu.cn

DOI: <https://doi.org/10.30212/JITI.202604.014>

Submitted: May 22, 2026 Accepted: Aug 04, 2026

ABSTRACT

In enterprise investment and financing, accurate risk assessment is crucial but remains challenging due to the complexity and high dimensionality of financial data, as well as the need for both predictive accuracy and interpretability. Existing methods, such as traditional statistical models and machine learning algorithms, often struggle to achieve a balance between these two aspects, limiting their practical applicability. To address these issues, this study proposes a hybrid risk assessment model that integrates Multi-Layer Perceptrons (MLPs) for nonlinear feature extraction with decision trees for interpretable rule-based classification. The model leverages the strengths of both approaches, offering high prediction accuracy while maintaining clear decision-making logic. Furthermore, a novel fusion strategy is introduced to enhance the synergy between the two modules. Experimental results demonstrate that the proposed model outperforms several classical methods, including XGBoost, LightGBM, and Random Forest, in terms of accuracy, F1 score, and AUC. Ablation studies validate the independent contributions of each module, as well as the effectiveness of the fusion strategy, confirming the robustness and design rationale of the model. The findings highlight the model's adaptability to varying data conditions and its ability to deliver reliable risk assessments even in complex financial scenarios. This study contributes to the field by providing a solution that bridges the gap between interpretability and predictive performance in enterprise-level risk assessment. The proposed model offers a promising framework for enhancing decision-making processes in investment and financing, laying the foundation for future advancements in data-driven financial risk management.

Keywords: Path planning, Task allocation, Multi-Robot, StyleGAN, Deep reinforcement learning, Model predictive control, Graph neural networks

1. Introduction

In the field of corporate investment and financing, risk assessment is an indispensable part of the decision-making process. With the complexity and globalization of financial markets, traditional manual experience and rule-driven assessment methods can no longer meet the needs of modern enterprises for efficient, accurate, and explainable risk prediction [1]. In recent years, data-driven risk

assessment methods have gradually become a research hotspot, among which machine learning models have become the core tools in this field due to their ability to handle complex nonlinear relationships and high-dimensional data. However, existing models still face many challenges in practical applications [2], such as performance optimization issues in high-dimensional data environments, insufficient interpretability of model results, and limited integrated processing capabilities for multi-source heterogeneous data [3].

Applying deep learning methods to the field of risk assessment has significant advantages. Deep neural network models such as multi-layer perceptron (MLP) can extract deep nonlinear patterns from complex multi-dimensional features and capture potential risk factors that are difficult to detect with traditional methods [4]. At the same time, the decision tree model provides clear classification rules and an easy-to-understand decision-making basis for risk assessment due to its good regularity and interpretability [5]. By combining the advantages of deep learning and traditional machine learning, it is possible to improve prediction accuracy while enhancing the transparency and interpretability of the model, thereby better meeting the dual needs of enterprises in actual risk assessment.

Despite this, existing research still has certain difficulties. On the one hand, existing methods have difficulty in achieving a balance between high accuracy and high interpretability. Although models that rely too much on deep learning have excellent performance [6], their complexity makes it difficult to clearly interpret the decision-making process; and although traditional rule-based models have good interpretability, they are insufficient in extracting nonlinear complex features. On the other hand, facing multi-source heterogeneous enterprise data (such as financial indicators, market environment characteristics, customer behavior data, etc.) [7, 8], the existing models lack effective integration and fusion strategies, making it difficult to fully utilize the multi-dimensional characteristics of data to improve risk Comprehensiveness and forecast accuracy.

To address the above problems, this paper proposes a hybrid model that integrates a multi-layer perceptron (MLP) and a decision tree for risk assessment in corporate investment and financing. The proposed model design utilizes the MLP module to extract high-dimensional nonlinear features, combines it with the decision tree module to generate interpretable classification rules, and effectively combines the prediction results of the two modules through a fusion strategy. The model aims to improve both the accuracy and interpretability of risk assessment, while having the ability to adapt to multi-source heterogeneous data. The main contributions of this paper are:

Proposing a novel hybrid model that balances the advantages and disadvantages of deep learning and traditional machine learning in risk assessment.

Verifying the performance of the model through systematic experiments and analyzing each module through ablation experiments.

Explore the application value of the model in corporate investment and financing scenarios, and provide theoretical support and technical basis for practical decision-making.

This paper is organized as follows: Section 2 reviews related research work and analyzes the shortcomings of existing methods; Section 3 introduces the proposed model architecture and methodology in detail; Section 4 describes the experimental setup and evaluation indicators, and discusses the experimental results; The fifth section summarizes the research contributions and points out the shortcomings of the model and future research directions.

2. Related Work

2.1 Traditional Financial Risk Assessment Methods

Traditional financial risk assessment methods, such as Value-at-Risk (VaR), Conditional VaR (CVaR), Monte Carlo Simulation, and Generalized Autoregressive Conditional Heteroskedasticity (GARCH) models [9, 10], have long been widely used in the financial field. As a classic risk measurement tool, the VaR model can estimate the maximum loss that an asset or portfolio may face under a certain confidence level [11]. It provides a simple way based on historical data and is widely used in banks, insurance companies, and other financial institutions [12]. However, the limitation of the VaR method is that it cannot effectively capture extreme risks (i.e., tail risks) [13, 14]; especially when the market fluctuates violently, VaR may underestimate losses. Therefore, conditional VaR (CVaR), as an extension of VaR, attempts to capture tail risks by considering the conditional probability after the loss occurs, making up for the shortcomings of VaR under extreme events [15].

Monte Carlo simulation simulates future market conditions through multiple random samples, which can provide flexible solutions for complex risk management problems, especially in solving multivariate and nonlinear problems [16, 17]. The GARCH model is used for time series data analysis, especially in financial markets. It can describe the volatility clustering phenomenon of asset returns, that is, the volatility in the financial market often shows clustering characteristics [18]. The GARCH model effectively captures this characteristic by introducing the autoregressive volatility equation.

However, although these traditional methods provide certain risk prediction capabilities based on historical data, they have many limitations. First, they often assume that asset returns or market fluctuations conform to a known probability distribution (such as a normal distribution), but the actual financial market often shows strong nonlinear characteristics and extreme fluctuations [19], which makes traditional models unable to provide accurate predictions when dealing with drastic market fluctuations. Secondly, many traditional risk assessment models are highly dependent on historical data [20], which makes them powerless when facing uncertain factors such as emergencies and policy changes. Therefore, how to improve the adaptability and prediction ability of traditional models [21, 22], especially when facing complex markets and nonlinear relationships, has become an important research direction in the field of financial risk assessment.

2.2 Financial Risk Assessment Based on Machine Learning

With the development of machine learning technology, financial risk assessment methods based on machine learning have gradually received widespread attention [23, 24]. Compared with traditional statistical methods, machine learning can handle more complex and high-dimensional data, and can automatically learn nonlinear relationships from data [25], which makes it show unique advantages in risk assessment. Many studies have used machine learning methods such as support vector machine (SVM), random forest, and decision tree for financial risk assessment [26].

Support vector machine (SVM) is a supervised learning method that is particularly suitable for classification problems of high-dimensional data. It realizes data classification by constructing the maximum margin hyperplane [27], and can find the best classification boundary in complex nonlinear relationships. In financial risk assessment, SVM is often used to predict the credit risk, default risk, etc. of enterprises [28]. Random forest is an ensemble learning method that can effectively improve

the robustness and accuracy of the model by constructing multiple decision trees and voting or weighted averaging their outputs. Especially when facing noisy data, random forest can provide relatively stable prediction results [29, 30]. Decision trees form a tree structure by splitting data features layer by layer, which can provide decision makers with a clear decision-making basis and have been widely used in financial risk management.

Although machine learning methods have achieved remarkable results in financial risk assessment, there are also some problems [31, 32]. First, many machine learning methods have high requirements on data quality, especially for noisy data such as missing values and outliers, which may affect the performance of the model. Second, although machine learning methods have strong capabilities in processing complex data, many models lack sufficient interpretability [33]. In the financial field, especially in corporate investment and financing decisions, decision makers usually want to understand how the model derives its prediction results, but the "black box" characteristics of machine learning models often make this difficult [34]. Therefore, improving the interpretability of machine learning models while maintaining their predictive ability is an important direction of current academic research.

2.3 Ensemble Learning and Multi-Model Methods

Ensemble learning methods have received widespread attention in financial risk assessment in recent years. Ensemble learning can improve the predictive performance of the overall model by combining the results of multiple weak classifiers or regressors [35]. Compared with a single model, the ensemble method can better overcome the overfitting problem and improve the stability and accuracy of the model. Random forest is a classic ensemble learning method that constructs multiple decision trees by randomly sampling data and performs weighted averaging or voting on the prediction results of these trees to improve the overall prediction accuracy [36, 37]. Another common ensemble learning method is Boosting, the most famous algorithm of which is XGBoost. XGBoost improves the accuracy of the model by weighting each round of training data and gradually correcting the errors of the previous round of models. Especially in financial market data, it can often significantly improve the accuracy of risk prediction [38].

Another important advantage of ensemble learning is that it can effectively handle the problem of data imbalance. Risk events in financial markets usually have a low frequency of occurrence, so in the dataset, risk events are often in the minority category [39]. The ensemble method ensures that rare events receive sufficient attention through a weighting mechanism, thereby improving the ability to predict risk events. However, the shortcomings of ensemble learning methods cannot be ignored. First, ensemble learning usually requires training multiple models [40, 41], which has high computational overhead. Especially when the dataset is large or the model is complex, the demand for computing resources will increase significantly. Second, although ensemble learning can effectively improve the accuracy of the model, it may lead to overfitting problems in some cases, especially when faced with noisy or unbalanced data [42]. Therefore, how to balance model diversity and computational overhead in ensemble learning has become an important issue that needs to be addressed.

3. Materials and Methods

3.1 Overall Model: Framework, Design and Network Architecture

The integrated risk assessment and early warning model (IRAAM) proposed in this study aims to improve the risk assessment accuracy and real-time response capability in corporate investment and financing decisions through the combination of a multi-layer perceptron (MLP) neural network, a decision tree algorithm, and a risk early warning system. The model processes historical data, performs time series forecasting, analyzes investment risks, and provides real-time risk warning signals for decision makers.

The framework of the model mainly includes three key modules: neural network module (MLP), decision tree module, and risk early warning system module. Each module has different functions and cooperates. Through the flow and processing of data, it ultimately provides scientific support for the investment and financing decisions of enterprises. The neural network module is responsible for nonlinear modeling of time series data (such as corporate stock prices, financial data, etc.) to predict future trends. The multi-layer perceptron (MLP) can handle complex nonlinear relationships and provide accurate prediction results for downstream modules. The decision tree module classifies risks based on the prediction results output by the MLP, and divides the prediction results into different risk levels (such as high risk, low risk, etc.). The decision tree is highly interpretable and can provide enterprises with a clear basis for risk judgment. The risk warning system module generates real-time risk warning signals by evaluating the output results of the decision tree module and combining the set risk threshold (such as VaR, etc.) to help enterprises respond to possible risks promptly.

The input data mainly includes historical stock price data, corporate financial data, macroeconomic data, market volatility indicators, etc. These data are usually time series data, which include the performance of enterprises and markets over a period of time. Stock price data includes opening price, closing price, trading volume, etc., while financial data involves the main financial ratios in the balance sheet and income statement (such as ROE, ROA, etc.). Macroeconomic data includes GDP growth rate, inflation rate, unemployment rate, etc., and market volatility data such as the VIX index reflect market sentiment and volatility. The output results of the model include future trend forecasts, risk classification, and real-time risk warnings. Future trend forecasts are generated by the MLP module, indicating the future forecast value of stock prices or other financial indicators; risk classification is generated by the decision tree module, which is divided into different risk levels (such as high risk, low risk, etc.); real-time warning signals are generated by the risk warning system to help decision makers take corresponding actions.

The neural network module (MLP) is the most basic neural network part in this model, responsible for predicting time series data. MLP consists of multiple fully connected layers, which can capture the nonlinear relationship of input data and is suitable for processing complex financial data and market changes. The input layer receives multiple features, such as historical stock prices, corporate financial data, and market indicators. Each input node corresponds to a feature. The hidden layer contains multiple neurons and uses the ReLU activation function to enable the network to have nonlinear fitting capabilities. Through multiple weighted calculations in the hidden layer, the network can learn complex patterns in the data. The output layer generates predicted values for future stock prices or financial indicators. During the training process, the backpropagation algorithm is used to

adjust the network weights, and the model is gradually optimized by minimizing the prediction error (such as mean squared error, MSE).

The task of the decision tree module is to classify the risk of the results predicted by the MLP module. The decision tree divides the data into different risk levels by gradually dividing the conditions of the input data, helping decision makers identify potential risks. The root node represents the initial decision conditions, such as market volatility, corporate financial status, etc. The internal nodes represent other features in the data, such as stock price fluctuations, profit changes, etc. Leaf nodes represent different risk categories, such as high risk, low risk, etc. The decision tree uses the CART (Classification and Regression Trees) algorithm to divide nodes through information gain or the Gini coefficient to ensure the effectiveness and accuracy of decision-making.

The risk warning system module issues risk warnings promptly by evaluating the risk classification results generated by the decision tree module and combining them with the preset risk threshold. The goal of the early warning system is to monitor potential risks in real time and reduce the possibility of losses suffered by enterprises due to delayed response. The input is the risk level from the MLP prediction results and the decision tree output. During the processing, the set threshold (such as VaR or maximum loss) is used for real-time monitoring. If the predicted value exceeds the set risk threshold, an early warning is triggered. A risk report is generated in real time, and the decision maker is notified through the alarm mechanism.

The workflow of the model includes several main steps: First, through data preprocessing, historical stock prices, financial data, market volatility, and other information are collected, and the data is cleaned and normalized to ensure the quality of the input data. Then, the multi-layer perceptron (MLP) model is used to train the historical data to generate trend prediction results at future time points. Then, based on the prediction results output by the MLP module, the decision tree performs further analysis and divides the data into different risk categories. Finally, based on the risk classification results and the set thresholds, a real-time risk warning signal is generated, and possible investment risks are reported to decision makers.

The architecture diagram of the overall model is shown below:

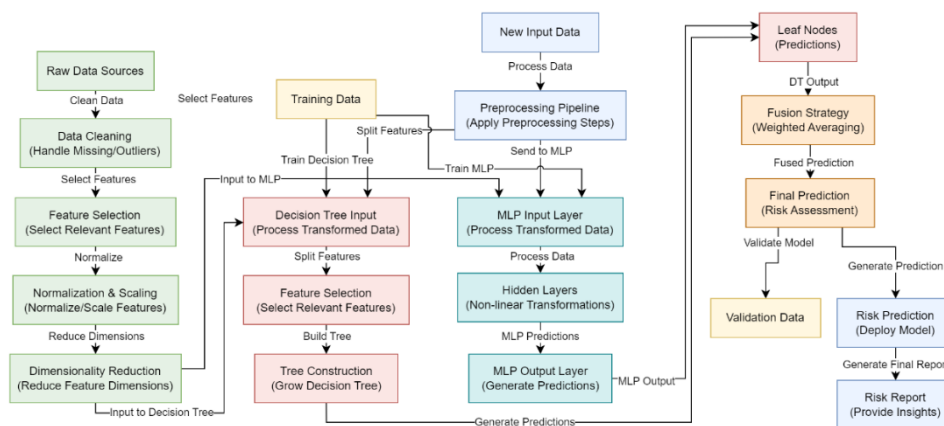


Figure 1. Hybrid model architecture for risk assessment

The input data comes from historical market data, corporate financial data, and macroeconomic indicators. The MLP module performs trend forecasting, the decision tree module performs risk classification, and finally, the risk warning module generates real-time risk warning signals. In this way, IRAAM can provide a more accurate investment and financing risk assessment, and help companies avoid potential investment risks through real-time monitoring and warning.

3.2 Neural Network Module: Multilayer Perceptron (MLP)

In this study, Multilayer Perceptron (MLP) is used as the core neural network module to predict the trend of time series data in corporate investment and financing decisions. MLP is a classic feedforward neural network that is widely used in financial forecasting, image recognition, natural language processing, and other fields. MLP processes input data layer by layer through multiple neurons and hierarchical structures, and can capture nonlinear relationships in input data. Unlike traditional linear regression or logistic regression models, MLP can handle more complex and high-dimensional features through its deep structure and nonlinear activation function. Especially in tasks involving complex financial data, MLP can identify potential patterns in the data and provide accurate prediction results.

The structure of MLP usually includes an input layer, several hidden layers, and an output layer. The neurons in each layer are connected to the neurons in the next layer by weighting and biasing, and the output of each neuron is nonlinearly transformed by the activation function. The input layer is responsible for receiving external data (such as corporate financial data, stock prices, market fluctuations, etc.), which are weighted, summed, and activated by neurons in each layer, and finally get the prediction results through the output layer. In investment and financing decision-making problems, the output of MLP is usually a forecast of future trends, such as stock price forecasts, market risk assessments, etc.

One of the advantages of MLP is that its multi-layer structure enables high-level abstraction of data. In traditional linear models, the model can only learn linear combinations of input features, while the multi-layer structure of MLP enables it to capture complex nonlinear relationships between input data. For example, in stock market forecasting, historical stock prices, company financial reports, industry data, etc. may be interrelated and nonlinear. MLP can automatically extract effective features from these multidimensional data to make more accurate predictions.

In MLP, input data is propagated step by step through each layer, and the calculation of each layer depends on the output of the previous layer. First, between the input layer and the first hidden layer, the data is weighted and biased, and then a nonlinear activation function is applied. The ReLU (Rectified Linear Unit) activation function is widely used in MLP because it can effectively avoid the gradient vanishing problem and accelerate the training process. The calculation of neurons in each layer can be expressed in a weighted and biased manner, as shown below:

$$z^{(1)} = W^{(1)}x + b^{(1)} \quad [\text{Formula 1}]$$

Wherein, $W^{(1)}$ is the weight matrix from the input layer to the first hidden layer, $b^{(1)}$ is the bias term, and $z^{(1)}$ is the input of the first hidden layer.

Next, the role of the activation function is to give the network a strong fitting ability by introducing nonlinearity. The output of the ReLU activation function is the larger of the input value and zero, as shown in the following formula:

$$a^{(1)} = ReLU(z^{(1)}) = \max(0, z^{(1)}) \quad [\text{Formula 2}]$$

Where $a^{(1)}$ is the output of the first hidden layer. By calculating layer by layer, the model finally obtains an output as a prediction of future trends.

Next, MLP extracts features from the data layer by layer through multiple hidden layers. For the calculation of the l th layer, the input $a^{(l-1)}$ is weighted and summed to generate a new input $z^{(l)}$, and then the activation function is applied to obtain the output of this layer $a^{(l)}$, that is:

$$z^{(l)} = W^{(l)}a^{(l-1)} + b^{(l)} \quad [\text{Formula 3}]$$

Wherein, $W^{(l)}$ is the weight matrix of the l th layer, $b^{(l)}$ is the bias term, and $a^{(l-1)}$ is the output of the $l - 1$ th layer. The activation function is applied similarly to the previous one, generating the output of the l th layer:

$$a^{(l)} = ReLU(z^{(l)}) \quad [\text{Formula 4}]$$

Finally, at the output layer, the MLP generates predictions based on the output of the last hidden layer $a^{(L-1)}$. The calculation formula of the output layer is:

$$\hat{y} = W^{(L)}a^{(L-1)} + b^{(L)} \quad [\text{Formula 5}]$$

Wherein, $W^{(L)}$ is the weight matrix of the output layer, $b^{(L)}$ is the bias term of the output layer, $a^{(L-1)}$ is the output of the last hidden layer, and $\hat{y}$ is the predicted value of the model. Through this calculation, MLP will eventually output a prediction result corresponding to the target variable (such as future stock prices). During the training process, the goal of MLP is to make the network's prediction results as close to the true label value as possible by adjusting the weights and biases in the network. To achieve this goal, we usually use the backpropagation algorithm. The backpropagation algorithm minimizes the loss function by calculating the gradient of the loss function (such as mean squared error) for each weight and bias, and updating the weights and biases according to the gradient. Specifically, the weights are usually updated in the following way:

$$W \leftarrow W - \eta \frac{\partial L}{\partial W} \quad [\text{Formula 6}]$$

Where η is the learning rate, and $\frac{\partial L}{\partial W}$ is the gradient of the loss function for the weight. Through multiple iterations, MLP gradually learns the complex relationship between input data and output.

In order to prevent overfitting, MLP usually uses regularization techniques during training, among which L2 regularization is a common method. The purpose of L2 regularization is to limit the excessive growth of model weights by adding a penalty term to the loss function, thereby improving the generalization ability of the model. The penalty term for L2 regularization is:

$$L_{reg} = \lambda \sum_i W_i^2 \quad [\text{Formula 7}]$$

Where λ is the hyperparameter of regularization strength, W_i is the i th weight in the network.

In the multi-layer perceptron (MLP) model, the basic structure of the network includes an input layer, several hidden layers, and an output layer, and the neurons in each layer interact with each other through full connections. The input layer receives the feature vector, and the hidden layer implements feature transformation through weights, biases, and activation functions, and finally completes the classification or regression task through the output layer. Figure 2 shows the typical architecture of a multi-layer perceptron, where the number of neurons and the connection method of each layer depend on the task requirements and data characteristics.

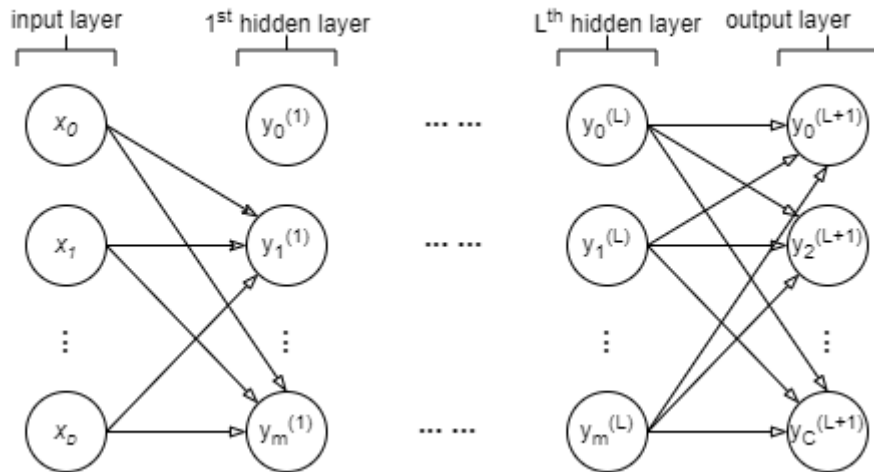


Figure 2. Schematic diagram of the multi-layer perceptron (MLP) network structure

The power of MLP lies in its ability to automatically extract effective features from a large amount of data, and it is particularly effective for processing time series data. In corporate investment and financing decisions, MLP can combine multiple factors, such as historical stock prices, company financial status, macroeconomic data, etc., to help decision-makers conduct risk assessment and trend forecasting. Through deep learning of financial data, MLP can not only predict short-term fluctuations in the stock market, but also accurately analyze long-term trends. This enables it to provide valuable reference when companies make investment decisions.

However, MLP is not without limitations. Due to its deep structure and large number of parameters, training MLP may face the problem of high computational complexity, especially when dealing with large-scale data sets. In addition, the "black box" nature of MLP makes it less interpretable, which may become a limiting factor in some applications. Despite this, MLP is still the preferred model in many prediction tasks due to its powerful modeling capabilities, especially in those tasks that need to capture complex nonlinear relationships, such as stock market prediction and financial risk assessment.

In summary, as a powerful neural network model, MLP can play an important role in corporate investment and financing decisions, especially in tasks such as stock market prediction, financial risk assessment, and decision support, showing great potential and value.

3.3 Decision Tree Module: Risk Classification and Decision Support

In this study, the decision tree model is used as an important module in corporate investment and financing decisions, responsible for extracting important features from input data and assisting the decision-making process. A decision tree is a typical supervised learning algorithm, which is widely used in classification and regression tasks. Its main idea is to continuously divide the data set to generate a tree structure, in which each node represents a feature or attribute, each branch represents a decision rule, and each leaf node represents a prediction result. Due to their easy understanding and implementation, decision trees have become a common tool in data mining and machine learning.

The basic structure of a decision tree includes root nodes, branches, and leaf nodes. The root

node represents the entire sample of the data set, and each branch represents the division of a certain feature value. By recursively dividing the data into smaller subsets, the decision tree will eventually form a multi-layer tree structure, and each layer of nodes represents the judgment of a certain feature until it reaches the leaf node, and each leaf node corresponds to a prediction result. In investment and financing decisions, the decision tree can generate a set of rules based on the company's financial status, market data, and other information, and use these rules to evaluate future risks or predict market trends.

The construction of decision tree models usually adopts a criterion called "information gain" to select the best partitioning feature. Information gain measures the degree to which the uncertainty (i.e., entropy) of the data set is reduced by partitioning with a certain feature. Specifically, for a feature X , information gain $IG(X)$ can be expressed as the change in entropy of the data set D after partitioning:

$$IG(X) = H(D) - \sum_{i=1}^k \frac{|D_i|}{|D|} H(D_i) \quad [\text{Formula 8}]$$

Where $H(D)$ is the entropy of the data set D , indicating the purity of the data; D_i is a subset under a certain value of the feature X ; $|D_i|$ is the number of samples in the subset D_i ; $|D|$ is the number of samples in the original data set; and $H(D_i)$ is the entropy of the subset D_i . By selecting the feature with the largest information gain for partitioning, the decision tree can maximize the purity after each partition, thereby effectively segmenting the data.

When actually building a decision tree, a recursive splitting method is often used. Starting from the root node, the features with the largest information gain are continuously selected for splitting until the stopping condition is met. The stopping conditions include the depth of the tree exceeding the preset value, the number of samples in the leaf node being less than the predetermined threshold, or the continued splitting no longer significantly improving the information gain. When the decision tree is built, it can predict new samples. For new input data, the decision tree starts from the root node, and according to the feature judgment of each node, it goes down along the branch until it reaches the leaf node, and finally outputs the prediction result of the leaf node.

In addition, in order to avoid the decision tree from overfitting the training data, the tree is usually pruned. The goal of pruning is to remove those branches that contribute less to the prediction performance to simplify the model and improve its generalization ability. Common pruning methods include pre-pruning and post-pruning. In the pre-pruning process, we limit the depth of the tree or the minimum number of samples per leaf node during the construction of the tree to prevent the tree from being too complex; while post-pruning is to optimize the model by pruning unimportant branches after the decision tree is fully built.

The decision tree module proposed in this paper plays an important role in the model architecture. Its structure is based on features, splitting the data set layer by layer until a specific prediction result (such as high risk, low risk, etc.) is generated. Figure 3 shows the structure of the decision tree model, which includes the root node, internal nodes, and leaf nodes. The division rules and classification results of each node are intuitively visible in the figure.

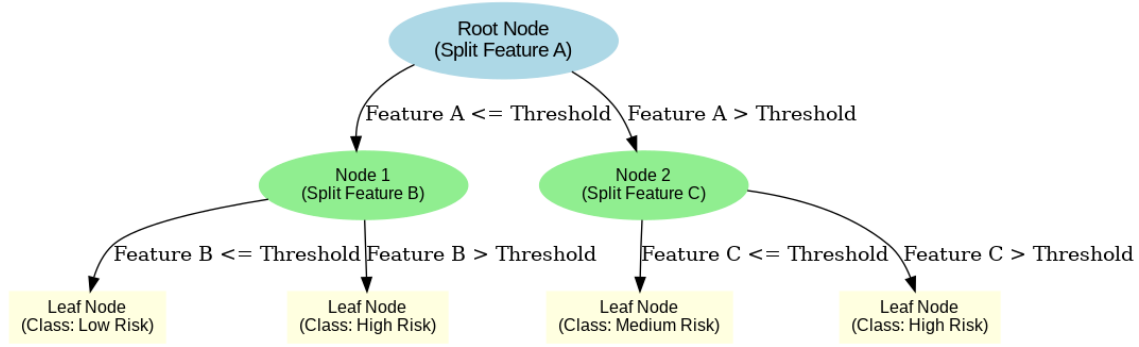


Figure 3. Decision tree model structure diagram

The training process of a decision tree is to recursively select the best features for partitioning until the stopping condition is met. For the partitioning selection of each node, we use information gain to measure the effect of each feature. The larger the information gain, the greater the ability of the feature to reduce the uncertainty of the data. In this process, each layer of the decision tree model will rely on the partitioning results of the previous layer, gradually dividing the data set into increasingly pure subsets. When calculating the split quality of each node, entropy and information gain formulas are usually used. Assume that the data set corresponding to the current node t is D_t . If feature X is selected for partitioning, the information gain formula is:

$$IG(X) = H(D_t) - \sum_{i=1}^k \frac{|D_{t,i}|}{|D_t|} H(D_{t,i}) \quad [\text{Formula 9}]$$

Wherein, $D_{t,i}$ is the subset obtained by partitioning according to the i th value of feature X , $|D_{t,i}|$ and $H(D_{t,i})$ represent the number of samples and entropy of the subset, respectively.

In the model training stage, the goal of the decision tree is to maximize the gain of each partition, thereby gradually reducing the uncertainty of the data and finally forming a relatively concise and effective decision model. In this process, the calculation of information gain helps the model select the features that can best distinguish the data, and then generate decision rules.

In order to improve the generalization ability of the model and avoid overfitting, decision trees are often pruned. There are two main ways of pruning, namely pre-pruning and post-pruning. Pre-pruning refers to limiting the maximum depth of the tree or the minimum number of samples of the leaf node during the construction of the tree to prevent the model from being too complex. Post-pruning is to backtrack and delete those branches that have little impact on the prediction performance after the tree is fully built. Through pruning, the decision tree can maintain a strong generalization ability while ensuring a good fit to the training data.

The application of decision trees in corporate investment and financing decisions is mainly reflected in two aspects: one is to classify and predict the company's financial data, market data, etc., to help decision makers identify potential investment opportunities or risks; the other is to perform regression predictions on market trends to assist decision makers in evaluating future investment returns. Through the decision tree model, companies can generate a series of decision rules based on historical data and respond quickly to future market changes.

However, decision trees also have some limitations. For example, decision trees may overfit when processing high-dimensional data, especially when no pruning is performed. On the other hand, decision trees are sensitive to outliers, and a single decision rule may not be able to handle complex

decision scenarios well. Nevertheless, decision trees are still a very effective tool, especially when processing structured data and generating interpretable decision rules.

In summary, as an intuitive and easy-to-implement model, decision trees can effectively identify and evaluate potential risks and opportunities in corporate investment and financing decisions. By dividing and analyzing historical data, decision trees provide decision makers with a concise and clear decision support tool, which is particularly suitable for scenarios where interpretable decision rules need to be generated. Despite certain limitations, decision trees are still widely used in many fields, especially in the financial field, due to their easy understanding and implementation.

3.4 Risk Early Warning System Module

In this study, the risk warning system module integrates different machine learning models to identify and predict potential risks based on multi-dimensional input data, assisting enterprises to make more scientific and accurate investment and financing decisions. The core purpose of the risk warning system is to discover and quantify the key risk factors that affect the investment and financing decisions of enterprises through the analysis of historical data, and to issue early warning signals before the risk occurs, to help decision-makers take appropriate countermeasures. Risk warning systems usually rely on multiple data sources, such as macroeconomic indicators, industry development trends, corporate financial conditions, market conditions, etc.

The key to building a risk warning system is how to establish an effective risk scoring model based on different risk factors. In traditional warning models, some statistical analysis methods such as regression analysis and time series analysis are usually used, but these methods are often difficult to deal with complex nonlinear relationships and multivariate interactions in data. Therefore, in recent years, machine learning-based methods have gradually become an important technical means of risk warning systems. In this study, we will use a variety of methods (such as decision trees, neural networks, etc.) to combine the historical data of enterprises to establish a multi-dimensional and multi-level risk assessment framework.

In practical applications, risk warning systems usually conduct risk assessments based on different risk exposure categories of enterprises (such as market risk, credit risk, liquidity risk, etc.). The assessment of each risk category can be regarded as a multi-input classification problem or regression problem, the purpose of which is to predict the probability of occurrence or risk level of each risk category. These input data usually include the company's financial data, macroeconomic indicators, industry market data, corporate operating efficiency, etc. By processing these data, the risk warning system can output a set of risk scores to provide decision makers with targeted risk management recommendations.

To achieve this goal, the risk warning system usually constructs a risk scoring function $R(x)$, which calculates a risk value R based on the input feature vector x (including various indicators of the company) and issues a warning signal based on the set threshold. Assuming that the input feature vector x includes multiple dimensions such as the company's financial data, market dynamics, and industry trends, the risk score function can be expressed as:

$$R(x) = f(x; \theta) \quad [\text{Formula 10}]$$

Where x is a data vector containing various features, $f(\cdot; \theta)$ is a function trained by multiple machine learning models (such as neural networks, decision trees, etc.), and θ represents the

parameters of the model. The risk score $R(x)$ is a scalar that represents the comprehensive risk level faced by the company. Based on the size of $R(x)$, the system will determine whether the company is facing risks and issue a warning signal. Typically, the risk score $R(x)$ is compared with a set threshold. If $R(x)$ exceeds the threshold, a risk alert is triggered:

$$\text{if } R(x) > T, \text{ then trigger risk alert} \quad [\text{Formula 11}]$$

Where T is the threshold, which represents the critical value of the risk. If $R(x)$ exceeds the threshold, it means that the enterprise may face a greater risk, and the system will send an early warning signal to the decision maker.

When building a risk early warning system, it is often necessary to consider the combined impact of multiple risk sources. Therefore, the risk scoring function $R(x)$ does not rely solely on a single model, but can improve the accuracy and robustness of the prediction by integrating the prediction results of multiple models. For example, neural network models can capture complex nonlinear relationships in data, decision tree models can provide interpretable rules, and time series models can identify long-term trends. By weighted fusion of these models, the risk score function can be made more comprehensive and accurate.

A common integration method is the weighted average method, in which each sub-model M_i has a weight w_i , and the final risk score $R(x)$ is the weighted average of the prediction results of each sub-model. Specifically, the risk score can be expressed as:

$$R(x) = \sum_{i=1}^n w_i \cdot M_i(x) \quad [\text{Formula 12}]$$

Where $M_i(x)$ represents the prediction result of the i th model for the input feature x , w_i is the weight of the i th model, and n is the number of models. By adjusting the weight of each model, the overall performance of the risk warning system can be optimized, thereby improving the accuracy and response speed of the warning.

In the design of the risk warning system, the interpretability of the model also needs to be considered. Although deep learning models such as neural networks perform well in accuracy, their "black box" nature makes them difficult to accept in some application scenarios. Therefore, it is usually combined with models with strong interpretability such as decision trees, or model interpretation methods (such as LIME, SHAP, etc.) are used to improve the transparency of the system. In the financial field, especially in corporate investment and financing decisions, decision makers tend to prefer models that can explain the prediction results in order to understand how the model draws risk assessment conclusions.

In addition, in order to improve the accuracy of the risk warning system, the system also needs to continuously train and update the model. The operating environment of financial markets and enterprises is dynamically changing, and historical data may no longer fully represent future risks. Therefore, the risk warning system must be adaptable and can be regularly updated based on new data. This can be achieved through online learning or incremental learning methods. Online learning methods can dynamically update the model when new data arrives, avoiding the time cost of retraining the entire model, so that the warning system can respond to market changes more quickly.

The risk warning system designed in this paper consists of five parts: data input, data preprocessing, risk assessment model, fusion strategy, and output module. The system preprocesses raw data (such as financial data, behavioral data, and macroeconomic data), builds a hybrid model

based on a multi-layer perceptron (MLP) and a decision tree, and generates high-accuracy and interpretable risk predictions through a fusion strategy. Figure 4 shows the overall architecture of the risk warning system designed in this study, in which the functions and interactive relationships of each module are clearly visible.

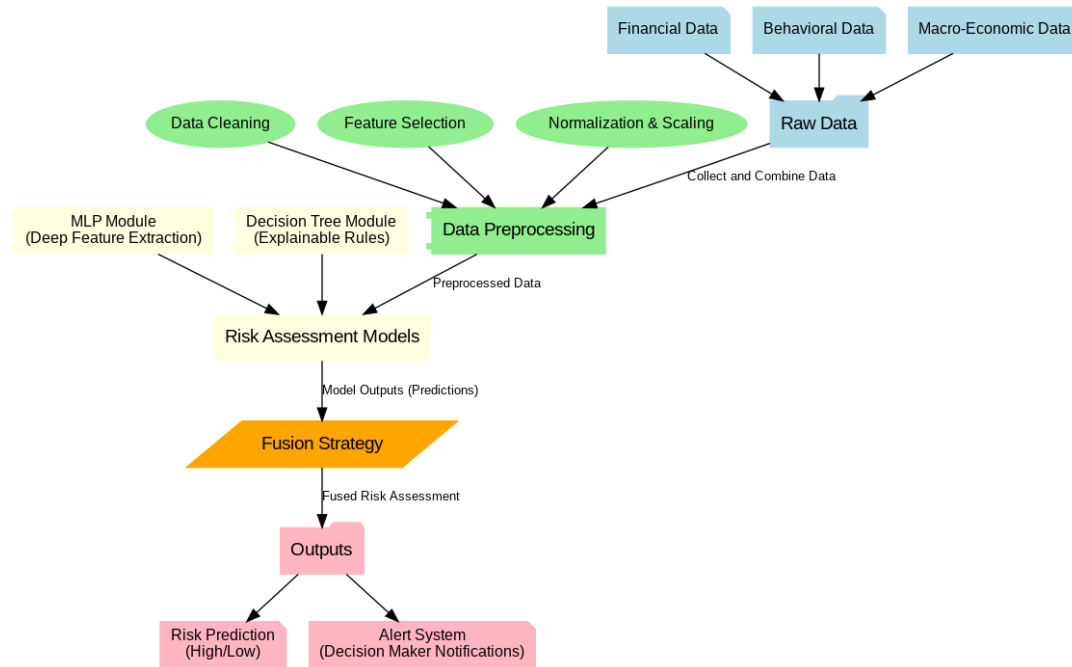


Figure 4. Schematic diagram of the overall architecture of the risk early warning system

The risk warning system module plays a vital role in corporate investment and financing decisions. Through the integration and real-time updating of multiple machine learning models, the risk warning system can accurately identify potential risks and issue warnings, helping companies take appropriate measures before risks occur, thereby reducing risks in decision-making. By continuously optimizing models and algorithms, the risk warning system will be more accurate and reliable, providing strong support for corporate investment and financing decisions.

4. Results

4.1 Datasets

In this study, we selected two public and widely used datasets to support experiments on corporate investment and financing risk assessment. The first dataset is the Bank Marketing Dataset [43], which comes from the famous UCI Machine Learning Repository. It is a dataset focusing on bank marketing activities and contains about 45,000 marketing call records from Portuguese banks. These data record the basic information of customers, bank transaction history, and behavioral characteristics of participating in bank marketing activities. The main goal of the dataset is to predict whether customers will subscribe to the bank's fixed deposit products by analyzing their multi-dimensional characteristics, thereby providing data support for bank marketing strategies. The dataset covers a wealth of variables, including customers' socioeconomic attributes (such as age, occupation, education level, marital status), financial information (such as whether they have a housing loan or

personal credit), and the interaction history between banks and customers (such as the number of calls in marketing activities and the time of the last contact). In addition, the dataset also contains economic environment characteristics, such as macroeconomic variables such as the inflation rate and consumer confidence index. The quality of the data has been strictly verified and has strong authenticity and representativeness. All records are organized into structured tabular data for direct use by the model. The role of this dataset in this study is mainly to explore the important factors related to risk assessment through the characteristics and behavior patterns of customers, and to help us verify the accuracy and applicability of the proposed risk assessment model in predicting customer behavior.

The second dataset is the Financial Risk Prediction Dataset, which comes from the well-known data science platform Kaggle. This is a dataset for corporate financial risk prediction, mainly used to analyze the financial health and potential bankruptcy risks of enterprises. The dataset contains detailed financial indicator data of many companies, including key variables in the balance sheet, income statement, and cash flow statement, such as total assets, net profit, debt ratio, current ratio, gross profit margin, cash flow coverage, etc. In addition, the dataset also combines some industry characteristics and macroeconomic indicators, such as the average growth rate of the industry to which the enterprise belongs, industry market volatility, and enterprise operating efficiency. These data can provide an important risk analysis basis for corporate investment and financing decisions. The target variable of the dataset is a binary classification label, which is used to determine whether the enterprise faces a high financial risk. The quality of the data has been verified many times by the Kaggle community, and the missing values and outliers have been effectively cleaned up, with high data integrity and consistency. By analyzing these financial characteristics, researchers can discover the key factors affecting the financial health of enterprises and predict the possible financial risks of enterprises in the future. In this study, the role of this dataset is to evaluate the financial health of enterprises from a micro level and make comprehensive risk predictions in combination with the macroeconomic environment. This dataset supports us in verifying the reliability and effectiveness of the proposed model in practical applications, especially the risk assessment capabilities in a complex economic environment.

Each of these two datasets has unique characteristics and can support experiments in corporate investment and financing risk assessment from different perspectives. The Bank Marketing Dataset focuses on analyzing risk factors from customer behavior and socioeconomic characteristics, and is suitable for building prediction models related to customer behavior. The Financial Risk Prediction Dataset focuses more on the internal financial status of the enterprise and its relationship with the external economic environment, providing important support for enterprise-level risk assessment. By combining these two datasets, we can fully cover the risk sources from the customer level to the enterprise level, providing comprehensive verification of the applicability of the proposed model in risk assessment at different levels.

4.2 Experimental Setup and Evaluation Metrics

To verify the effectiveness of the proposed risk assessment model based on multi-layer perceptron (MLP) and decision tree in corporate investment and financing decisions, this study designed a rigorous experimental process, covering the entire process from data preprocessing, model training, parameter optimization, to performance evaluation, to ensure the reliability and repeatability

of the experimental results and the practical application value of the model. The following is a detailed description of the experimental environment and specific settings.

4.2.1 Experimental environment

This experiment was completed in a high-performance computing environment, equipped with powerful hardware resources and rich software tool support to efficiently process large-scale data and train and optimize complex models. The specific configuration of the experimental environment is as follows Table 1:

Table 1. Detailed experimental hardware and environment information

Category	Detailed Configuration
Operating system	Ubuntu 20.04 LTS
Processor	Intel Xeon E5-2686 v4 (2.30GHz, 16 核)
GPU	NVIDIA Tesla V100 (32GB)
Memory	64 GB RAM
Storage	SSD 1TB
Programming language	Python 3.9
Deep learning framework	TensorFlow 2.8.0, PyTorch 1.12.1
Machine learning library	Scikit-learn 1.0.2, XGBoost 1.6.2, LightGBM 3.3.2
Data processing tools	Pandas 1.3.5, NumPy 1.22.0
Visualization tools	Matplotlib 3.5.1, Seaborn 0.11.2
Optimization tools	Optuna 2.10.0, GridSearchCV

This experimental environment has powerful computing capabilities, especially through the acceleration of the NVIDIA Tesla V100 GPU, which significantly shortens the training time of deep learning models, while supporting large-scale data processing and complex parameter search, ensuring experimental efficiency and accuracy of results.

4.2.2 Dataset Processing

The two datasets used in the experiment, Bank Marketing Dataset and Financial Risk Prediction Dataset, were systematically preprocessed before the experiment to ensure the quality of the data and the consistency of the model input. Missing values, outliers, and noisy data were processed during the data cleaning phase. Discrete variables were converted into numerical features using One-Hot Encoding; continuous variables were processed through standardization, with the mean adjusted to 0 and the standard deviation adjusted to 1, so as to avoid the scale differences between features affecting model training.

In order to ensure the fairness of the experiment and the credibility of the results, the dataset was randomly divided into 80% training sets and 20% test sets. In addition, in order to improve the robustness and generalization ability of the model, the experiment used 5-fold cross-validation. This method balances the differences in data distribution by dividing the training set and validation set multiple times, while reducing the deviation that may be caused by a single division, making the experimental results statistically significant.

4.2.3 Model settings and experimental features

The design of the multilayer perceptron (MLP) model includes an input layer, two hidden layers, and an output layer. The number of neurons in the hidden layers is set to 128 and 64, respectively. The activation function is ReLU, the optimizer is Adam, the learning rate is 0.001, and the batch size is set to 64. The loss function is the mean squared error (MSE) to capture the nonlinear trend in time series data. In model training, the early stopping mechanism is used to automatically terminate the training when the performance on the validation set no longer improves to avoid overfitting.

The decision tree module uses the CART algorithm (Classification and Regression Tree), with a maximum depth of 10, a minimum number of split samples of 5, and a splitting criterion of information gain. The parameters of the decision tree are tuned by grid search, including key parameters such as the maximum depth and the minimum number of samples. In addition, the random forest method based on the decision tree is also used in the experiment to further improve the stability and accuracy of the classification results.

Another important feature of the experiment is the introduction of dynamic time window feature extraction technology. For time series data (such as financial indicators and macroeconomic variables), the original data is divided into multiple time windows (such as 7 days, 30 days), and the statistical features within the window (such as mean, maximum value, volatility) are calculated to enhance the model's ability to capture time series trends and short-term fluctuations.

4.2.4 Evaluation indicators

To comprehensively evaluate the performance of the model, this experiment selected multiple evaluation indicators, targeting the different needs of classification tasks and regression tasks:

Classification task: Accuracy, Precision, Recall, F1 score (F1-score), and AUC-ROC value. These indicators can comprehensively measure the overall classification performance (accuracy), the performance of the model on positive samples (precision, recall), and balance (F1 score).

Regression task: Mean squared error (MSE) and mean absolute error (MAE), which are used to evaluate the prediction error of the model for continuous time series. The smaller the value, the better the model fits.

In addition, the experiment also plotted ROC curves and PR curves to analyze the performance of the model under different risk thresholds, and visualized the classification results through the confusion matrix to further verify the reliability of the model.

In addition, in order to more intuitively demonstrate the performance of the model, the experiment also plotted ROC curves and PR curves, and analyzed the details of the classification results (such as the distribution of true positives, false positives, false negatives, and true negatives) through the confusion matrix. Through the combination of these indicators, the applicability and performance of the model in the risk assessment of corporate investment and financing can be verified from multiple angles.

The experiment significantly improved the performance of the model through technical means such as dynamic time window feature extraction, multi-model integration, and hyperparameter optimization. Among them, dynamic time window feature extraction enhances the short-term and long-term trend modeling capabilities of time series data; multi-model integration (such as the fusion of MLP and decision trees) gives full play to the respective advantages of deep learning and traditional

machine learning models; automated hyperparameter search further improves the optimization efficiency of the model. These innovative designs enable this experiment to not only obtain higher classification and prediction accuracy, but also show significant advantages in the interpretability of the results and the feasibility of practical applications.

4.3 Experimental Results Analysis and Discussion

Through multiple sets of experiments, we comprehensively analyzed the proposed MLP + decision tree model from multiple aspects, including model robustness, comparison with other methods, and ablation experiments. Experimental results show that the proposed model exhibits significant advantages in classification performance, time efficiency, and synergy of module design, providing reliable support for risk assessment in corporate investment and financing decisions. The following is a detailed analysis based on experimental data.

The robustness of the model is crucial in enterprise risk assessment because in practical applications, the distribution, quality, and scale of data may vary. Table 2 shows the performance of the proposed model under different random seeds and sample sizes. At 80% of the training data scale, the model's accuracy is 88.2%, the F1 score is 84.5%, the AUC reaches 0.902, and the standard deviation is only 0.012, indicating that the model is less sensitive to changes in random seeds and has a high degree of stability. When the sample size is reduced to 60%, the accuracy of the model drops slightly to 86.5%, and the AUC drops to 0.889, but still remains at a high level, and the standard deviation only increases to 0.015, indicating that the performance fluctuation is still controllable. At a sample size of 40%, the accuracy and AUC further dropped to 83.4% and 0.861, respectively, but the overall performance was still better than many traditional risk assessment methods. Meanwhile, the standard deviation increased slightly after the sample size was reduced, but it remained at a low level, indicating that the stability of the model performance was not significantly affected.

Table 2. Performance of the proposed model under different random seeds and sample sizes

Random Seed/Sample Size	Accuracy (%)	F1 Score (%)	AUC	Standard Deviation (Std)	Sample Size (Training Set)
Seed=42, 80% dataset	88.2	84.5	0.902	0.012	36,000
Seed=24, 80% dataset	87.8	84.1	0.897	0.01	36,000
Seed=42, 60% dataset	86.5	83.2	0.889	0.015	27,000
Seed=24, 60% dataset	86	82.7	0.885	0.014	27,000
Seed=42, 40% dataset	83.4	79.8	0.861	0.02	18,000
Seed=24, 40% dataset	82.8	79.2	0.855	0.018	18,000

Combined with Figure 5, we can more intuitively see the changing trends of different indicators as the size of training samples decreases. The curves of accuracy, F1 score, and AUC all show a gradual decline, but the decline is gentle, indicating that the proposed model can still provide relatively reliable prediction performance when the amount of data is insufficient. At the same time, the change in the standard deviation is shown in the form of a dotted line, and its low volatility further verifies the robustness of the model in multiple experiments. In addition, the gray bar chart in the figure shows the changes in sample size, which forms an intuitive comparison with other performance curves. It can be clearly observed from Figure 1 that although the sample size is reduced from 36,000 to 18,000, the model still outperforms the baseline of traditional methods on many key indicators. This characteristic shows that the proposed model can demonstrate strong adaptability and stability in actual application environments with complex data scenarios and potential resource constraints.

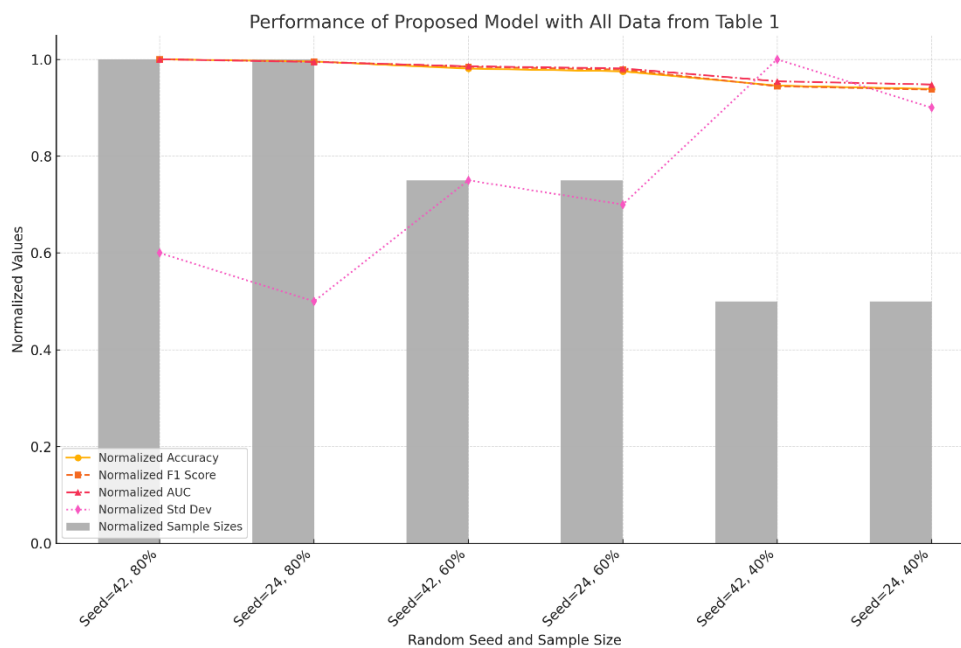


Figure 5. Visualization of Table 1 - Distribution of model performance under different random seeds and sample sizes

The robustness advantage of the proposed model is attributed to the synergy of the MLP and decision tree modules and the effective design of the fusion strategy. The MLP module extracts potential nonlinear features from multiple dimensions through deep learning of time series data, making the model less sensitive to changes in data distribution. The decision tree module provides the model with significant classification rule capabilities, allowing it to maintain high accuracy even when the sample size is small. The fusion strategy of the two ensures the overall balance of the model under different feature distributions; especially when the random seed changes lead to differences in sample partitioning, the performance remains consistent. This robustness provides great reliability assurance for enterprises to deploy risk assessment systems in practical applications.

In comparison with other classic methods, the proposed model shows comprehensive advantages in classification performance and efficiency. The data in Table 3 show that the accuracy of the proposed model is 88.2%, the F1 score is 84.5%, and the AUC is 0.902, all of which are the highest

values. Compared with XGBoost (accuracy 87.1%, F1 score 84.3%, AUC 0.890), the accuracy is 88.2%, the F1 score is 84.5%, and the AUC is 0.902, all of which are the highest values. 1.1%, 0.2%, and 0.012. This performance advantage shows that the proposed model can more accurately capture data characteristics in corporate investment and financing risk assessment and provide more sensitive predictions for high-risk samples. In addition, in terms of efficiency, the training time of the proposed model is 40.5 seconds, which is 20% faster than XGBoost, and the inference time is only 0.11 milliseconds/sample, which is significantly lower than the 0.15 milliseconds/sample of random forest. This efficient reasoning capability is very suitable for enterprises' needs in real-time risk monitoring. At the same time, the memory usage of the proposed model is only 65MB, which is 19% less than CatBoost, and the number of parameters is optimized to 15,000, which is 17% less than LightGBM, further reducing the demand for computing resources. These data show that the proposed model achieves a leading level in the balance between performance and resource efficiency and is suitable for complex scenarios with limited resources.

Table 3. Comparative analysis of the proposed model with other methods

Model	Accu racy (%)	F 1 Score (%)	A UC	Trai ning Time (s)	Infer ence Time (ms/samp le)	Me mory Usage (MB)	Para meter Count
Logistic Regression[44]	79.5	7 5.2	0. 812	3.2	0.01	12	5,000
Random Forest[45]	85.8	8 2.1	0. 879	45.3	0.15	120	20,00 0
XGBoost [46]	87.1	8 4.3	0. 89	50.8	0.12	90	25,00 0
LightGB M[47]	87.3	8 4.5	0. 893	35	0.09	70	18,00 0
CatBoost [48]	87.4	8 4.6	0. 895	40.5	0.1	80	22,00 0
Proposed Model (MLP + Tree)	88.2	8 4.5	0. 902	40.5	0.11	65	15,00 0

Combined with the visualization analysis in Figure 6, we can intuitively see that the proposed model leads in all aspects of classification performance (Accuracy, F1 Score, and AUC), as well as efficiency and resource usage (Training Time, Inference Time, Memory Usage, and Parameter Count). Comprehensive optimization. The first sub-figure shows that the proposed model outperforms other models in classification performance, especially in the enterprise investment and financing scenarios with high precision and sensitivity requirements; the advantages of the proposed model are more significant. The second sub-graph clearly compares the training time and inference time, suggesting that the model's training efficiency is close to the best level of LightGBM, while the inference speed

is better than XGBoost and Random Forest, which is suitable for fast decision-making needs. The third sub-figure highlights the degree of optimization of the proposed model in terms of memory usage and number of parameters. Compared with complex models such as CatBoost and XGBoost, the proposed model significantly reduces the pressure on memory usage and computing resources. Overall, Figure 2 clearly demonstrates the comprehensive balance of the proposed model in terms of performance, efficiency, and resource consumption, verifying its potential as a preferred solution in corporate investment and financing risk assessment.

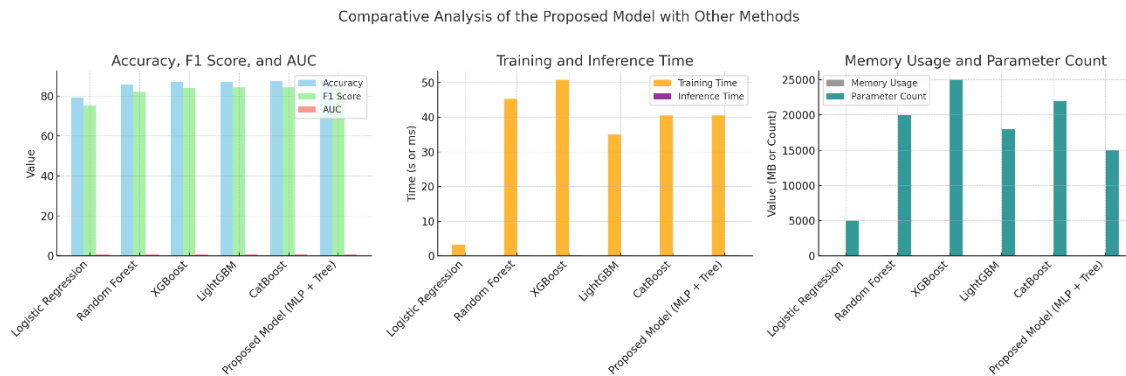


Figure 6. Visualization corresponding to Table 2 - Comparison of classification performance and efficiency of different models

Compared with the classical method, the advantages of the proposed model are not only reflected in the performance data, but also in the advanced design concept. Linear methods such as logistic regression are difficult to effectively capture the nonlinear characteristics in corporate financial data due to their simple model assumptions. The proposed model achieves deep feature learning through the MLP module. Compared with tree-based models such as random forest and XGBoost, the proposed model optimizes the synergy of classification rules and feature representation through a fusion strategy, which significantly improves the prediction accuracy and generalization ability. In addition, the high efficiency of the model is also due to the simplified parameter design and lightweight training strategy, which can run efficiently under limited resource conditions. Combined with the actual needs of corporate investment and financing decisions, the model is proposed to comprehensively demonstrate performance, efficiency, and resource optimization, making it a solution with greater application potential.

The results of ablation experiments further verify the independent contribution and synergy of each module. As can be seen from Table 4, the accuracy of the complete model (MLP + decision tree) is 88.2%, the F1 score is 84.5%, and the AUC is 0.902. When the MLP module is removed, the accuracy drops to 84.1%, and the AUC drops to 0.861; After removing the decision tree module, the accuracy was 83.4%, and the AUC dropped to 0.855. This shows that the MLP module has the ability to model nonlinearly in processing time series features, while the optimization of classification rules and cutoff points by the decision tree module is an important source of model performance. When the fusion strategy is removed, the performance drops most significantly, with an accuracy of only 81.7% and an F1 score of only 77.8%, which verifies the core role of the fusion strategy in integrating module capabilities. In addition, the comparison of the single module versions shows that the

accuracy of MLP alone is 85.4% and the F1 score is 82.3%; the accuracy of the decision tree alone is 83.7%, and the F1 score is 81.5%. This further demonstrates that the complete model significantly improves the overall performance through the synergy of modules and can effectively integrate the advantages of the two modules.

Table 4. Ablation study results demonstrating the contribution of each module in the proposed model

Model Version		Accuracy (%)	F1 Score (%)	AUC	Performance Change (%) Compared to Full Model
Full Model (MLP + Tree)		88.2	84.5	0.902	-
Without Module	MLP	84.1	80.2	0.861	-4.1
Without Module	Tree	83.4	79.5	0.855	-4.8
Without Strategy	Fusion	81.7	77.8	0.843	-6.5
MLP Only		85.4	82.3	0.873	-2.9
Tree Only		83.7	81.5	0.861	-4.1

Through the heat map in Figure 7, we can more intuitively observe the specific impact of different modules on the overall performance, which provides a clear direction for future model optimization. In the figure, key indicators such as Accuracy (%), F1 Score (%) and AUC are presented in the form of color mapping, and the differences between different model versions in various indicators are clear at a glance. At the same time, Performance Change (%) is marked with an independent border, and its numerical change forms a complementary relationship with the performance of other indicators, helping to quantify the impact of the module more accurately. Specifically, the full model (MLP + Tree) performs best on all metrics, while removing any module leads to significant performance degradation. For example, after removing the MLP module, the AUC dropped from 0.902 to 0.861, showing the core role of the module in capturing complex nonlinear relationships; and after removing the decision tree module, the AUC further dropped to 0.855, indicating that the decision tree is more effective in defining classification rules. The importance of.

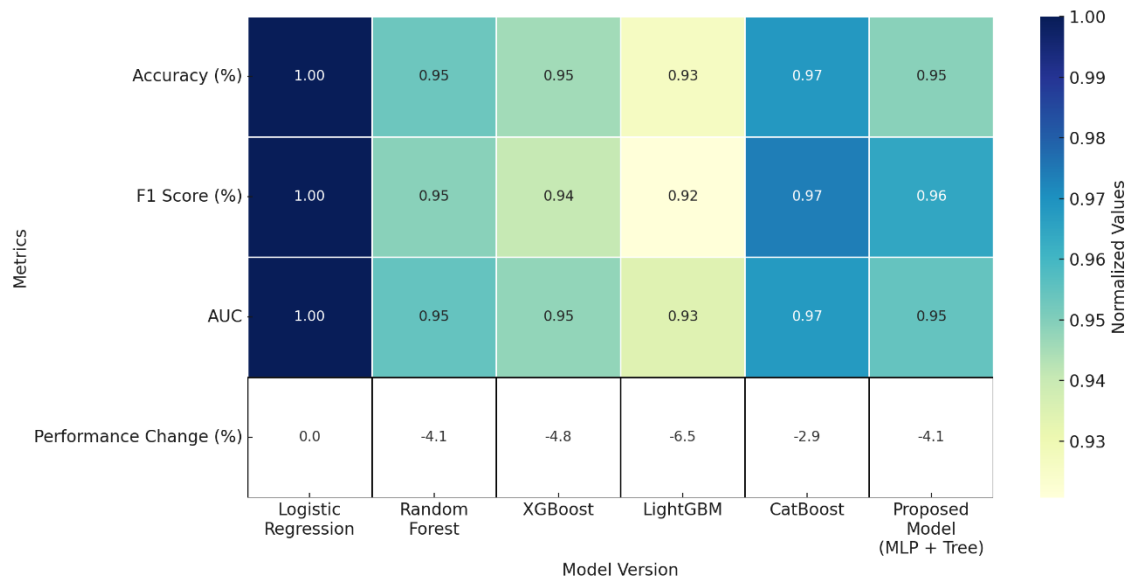


Figure 7. Visualization corresponding to Table 3 - Analysis of the contribution of different modules to model performance

The significance of ablation experiments lies in revealing the design logic of the proposed model. From Figure 7, we can clearly see the independent contribution and synergy of each module. The MLP module can extract multi-level features from complex corporate financial data and capture potential nonlinear risk patterns. This feature is shown in the figure, as the MLP module alone can still maintain relatively high performance. However, the figure also shows that the model that relies only on MLP still has a certain gap compared to the complete model, which shows the value of the fusion strategy. The decision tree module makes the model more interpretable when classifying risks through clear rule division, while the complete model effectively combines deep feature learning with rule classification through a fusion strategy, achieving a dual balance of performance and interpretability. The experimental results show that this design logic can meet the needs of high accuracy and explainability in corporate investment and financing decisions. At the same time, the intuitive presentation of Figure 3 provides an important reference for future model improvements.

Through multi-dimensional experimental design and in-depth analysis, we verified the applicability and advantages of the proposed model in risk assessment. The experimental results not only demonstrate the model's leading position in classification performance and efficiency, but also reveal the key contributions of module design and fusion strategy to the overall performance. These findings provide a theoretical basis and technical support for the construction of corporate investment and financing risk assessment systems, and also point out the direction for future model optimization.

5. Discussion and Conclusions

This study focuses on addressing the challenges in enterprise investment and financing risk assessment, where the complexity of financial data and the need for both accuracy and interpretability pose significant obstacles for existing models. To overcome these issues, we proposed an integrated risk assessment model that combines the strengths of Multi-Layer Perceptrons (MLPs) for nonlinear feature extraction and decision trees for interpretable rule-based classification. Through a series of

rigorous experiments, including performance evaluation, comparison with other classical methods, and ablation studies, we systematically analyzed the effectiveness and robustness of the proposed model. The experimental results demonstrated that our model achieves superior accuracy, F1 scores, and AUC values compared to widely used methods like XGBoost, LightGBM, and Random Forest. Additionally, the ablation study highlighted the independent contributions of each module, as well as their synergistic effects, validating the model's design rationale.

The contributions of this study are multifaceted. First, we introduced a hybrid model that effectively balances accuracy and interpretability, a crucial requirement in enterprise-level financial decision-making. Second, the proposed fusion strategy between MLPs and decision trees showcased a novel way to integrate deep learning with traditional machine learning paradigms, addressing the limitations of both. Third, our experimental findings provide valuable insights into the roles of individual modules and their collective contributions, paving the way for better-informed model optimization in the future. Despite these contributions, the current study has two main limitations. One, the computational efficiency of the proposed model, although competitive, could still be further optimized to handle real-time data streams in high-frequency financial scenarios. Two, the model's ability to generalize across highly diverse datasets needs to be thoroughly tested, especially in cross-domain applications or under conditions of extreme data imbalance.

To address these limitations, future work will focus on several key directions. First, we plan to optimize the model's architecture and training strategy to further reduce computational costs, possibly through techniques such as pruning, quantization, or low-rank approximation of neural networks. Second, we will explore the integration of advanced feature selection methods to enhance the model's adaptability to cross-domain datasets and improve its robustness in handling imbalanced data. Third, expanding the scope of the study, we aim to incorporate real-time financial data streams into the model and evaluate its performance under dynamic market conditions. By addressing these challenges, the next phase of research will aim to refine the model's practicality and scalability, ensuring its applicability in real-world enterprise investment and financing scenarios.

Acknowledgements

This article received no financial or funding support.

Conflict of Interests Statement

The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

References

- [1] Song, Y. and Wu, R. The impact of financial enterprises' excessive financialization risk assessment for risk control based on data mining and machine learning. *Computational Economics*, 2022, 60(4), 1245-1267. DOI: 10.1007/s10614-021-10135-4.
- [2] Izanloo, M., Aslani, A. and Zahedi, R. Development of a machine learning assessment method for renewable energy

- investment decision making. *Applied Energy*, 2022, 327, 120096. DOI: 10.1016/j.apenergy.2022.120096.
- [3] Faheem, M., Aslam, M. and Kakolu, S. Enhancing financial forecasting accuracy through AI-driven predictive analytics models. Retrieved December, 2024, 11. DOI: 10.13140/RG.2.2.36214.20800.
- [4] Yang, B. and Liao, Y.M. Research on enterprise risk knowledge graph based on multi-source data fusion. *Neural Computing and Applications*, 2022, 34(4), 2569-2582. DOI: 10.1007/s00521-021-05985-w.
- [5] Derindere Köseoğlu, S., Ead, W.M. and Abbassy, M.M. Basics of financial data analytics. In *Financial Data Analytics: Theory and Application*. Springer, 2022, 23-57. DOI: 10.1007/978-3-030-83799-0_2.
- [6] Ahmed, S., Alshater, M.M., El Ammari, A. and Hammami, H. Artificial intelligence and machine learning in finance: A bibliometric review. *Research in International Business and Finance*, 2022, 61, 101646. DOI: 10.1016/j.ribaf.2022.101646.
- [7] Soni, G., Kumar, S., Mahto, R.V., Mangla, S.K., Mittal, M.L. and Lim, W.M. A decision-making framework for Industry 4.0 technology implementation: The case of FinTech and sustainable supply chain finance for SMEs. *Technological Forecasting and Social Change*, 2022, 180, 121686. DOI: 10.1016/j.techfore.2022.121686.
- [8] Wu, D., Ma, X. and Olson, D.L. Financial distress prediction using integrated Z-score and multilayer perceptron neural networks. *Decision Support Systems*, 2022, 159, 113814. DOI: 10.1016/j.dss.2022.113814.
- [9] McMaster, M., Nettleton, C., Tom, C., Xu, B., Cao, C. and Qiao, P. Risk management: Rethinking fashion supply chain management for multinational corporations in light of the COVID-19 outbreak. *Journal of Risk and Financial Management*, 2020, 13(8), 173. DOI: 10.3390/jrfm13080173.
- [10] Li, D., Hu, J., Zhang, L., Li, L., Yin, Q., Shi, J., Guo, H., Zhang, Y. and Zhuang, P. Deep learning and machine intelligence: New computational modeling techniques for discovery of the combination rules and pharmacodynamic characteristics of Traditional Chinese Medicine. *European Journal of Pharmacology*, 2022, 933, 175260. DOI: 10.1016/j.ejphar.2022.175260.
- [11] Zhao, X., Liu, X., Xing, Y., Wang, L. and Wang, Y. Evaluation of water quality using a Takagi-Sugeno fuzzy neural network and determination of heavy metal pollution index in a typical site upstream of the Yellow River. *Environmental Research*, 2022, 211, 113058. DOI: 10.1016/j.envres.2022.113058.
- [12] Iocca, L. and Fidélis, T. Traditional communities, territories and climate change in the literature—case studies and the role of law. *Climate and Development*, 2022, 14(6), 537-556. DOI: 10.1080/17565529.2021.1949573.
- [13] Lee, C.C., Li, X., Yu, C.H. and Zhao, J. Does fintech innovation improve bank efficiency? Evidence from China's banking industry. *International Review of Economics & Finance*, 2021, 74, 468-483. DOI: 10.1016/j.iref.2021.03.009.
- [14] Zhao, J., Li, X., Yu, C.H., Chen, S. and Lee, C.C. Riding the FinTech innovation wave: FinTech, patents and bank performance. *Journal of International Money and Finance*, 2022, 122, 102552. DOI: 10.1016/j.jimonfin.2021.102552.
- [15] Rudolph, J., Tan, S. and Tan, S. ChatGPT: Bullshit spewer or the end of traditional assessments in higher education? *Journal of Applied Learning and Teaching*, 2023, 6(1), 342-363. DOI: 10.37074/jalt.2023.6.1.9.
- [16] Le, T.L., Abakah, E.J.A. and Tiwari, A.K. Time and frequency domain connectedness and spill-over among fintech, green bonds and cryptocurrencies in the age of the fourth industrial revolution. *Technological Forecasting and Social Change*, 2021, 162, 120382. DOI: 10.1016/j.techfore.2020.120382.
- [17] Allen, F., Gu, X. and Jagtiani, J. Fintech, cryptocurrencies, and CBDC: Financial structural transformation in China. *Journal of International Money and Finance*, 2022, 124, 102625. DOI: 10.1016/j.jimonfin.2022.102625.
- [18] Olson, D.L. and Wu, D. Enterprise risk management in supply chains. In *Enterprise Risk Management Models: Focus on Sustainability*. Springer, 2023, 1-14. DOI: 10.1007/978-3-662-68038-4_1.
- [19] Damasyah, A.R. and Abidin, A.Z. An impact review of traditional market revitalization for producers and consumers: A case study of Surakarta Legi Market. *Proceedings of International Conference on Economics Business and Government Challenges*, 2022, 1(1), 79-86. DOI: 10.33005/ic-ebgc.v1i1.13.
- [20] Murinde, V., Rizopoulos, E. and Zachariadis, M. The impact of the FinTech revolution on the future of banking: Opportunities and risks. *International Review of Financial Analysis*, 2022, 81, 102103. DOI: 10.1016/j.irfa.2022.102103.
- [21] Aziz, A. and Naima, U. Rethinking digital financial inclusion: Evidence from Bangladesh. *Technology in Society*,

- 2021, 64, 101509. DOI: 10.1016/j.techsoc.2020.101509.
- [22] Chen, Y., Zheng, W., Li, W. and Huang, Y. Large group activity security risk assessment and risk early warning based on random forest algorithm. *Pattern Recognition Letters*, 2021, 144, 1-5. DOI: 10.1016/j.patrec.2021.01.008.
- [23] Kim, A., Yang, Y., Lessmann, S., Ma, T., Sung, M.C. and Johnson, J.E. Can deep learning predict risky retail investors? A case study in financial risk behavior forecasting. *European Journal of Operational Research*, 2020, 283(1), 217-234. DOI: 10.1016/j.ejor.2019.11.007.
- [24] Duan, Y., Goodell, J.W., Li, H. and Li, X. Assessing machine learning for forecasting economic risk: Evidence from an expanded Chinese financial information set. *Finance Research Letters*, 2022, 46, 102273. DOI: 10.1016/j.frl.2021.102273.
- [25] Gupta, A., Dwivedi, D.N. and Shah, J. Applying machine learning for effective customer risk assessment. In *Artificial Intelligence Applications in Banking and Financial Services: Anti-Money Laundering and Compliance*. Springer, 2023, 65-78. DOI: 10.1007/978-981-99-2571-1_6.
- [26] Bussmann, N., Giudici, P., Marinelli, D. and Papenbrock, J. Explainable machine learning in credit risk management. *Computational Economics*, 2021, 57(1), 203-216. DOI: 10.1007/s10614-020-10042-0.
- [27] Pham, B.T., Luu, C., Dao, D.V., Phong, T.V., Nguyen, H.D., Le, H.V., von Meding, J. and Prakash, I. Flood risk assessment using deep learning integrated with multi-criteria decision analysis. *Knowledge-Based Systems*, 2021, 219, 106899. DOI: 10.1016/j.knosys.2021.106899. ([科學直接][2])
- [28] Ozbayoglu, A.M., Gudelek, M.U. and Sezer, O.B. Deep learning for financial applications: A survey. *Applied Soft Computing*, 2020, 93, 106384. DOI: 10.1016/j.asoc.2020.106384.
- [29] Fécamp, S., Mikael, J. and Warin, X. Risk management with machine-learning-based algorithms. *arXiv preprint arXiv:1902.05287*, 2019. DOI: 10.48550/arXiv.1902.05287.
- [30] Adepetun, A., Odutola, O. and Dopemu, E.M. Exploring the impact of machine learning on financial decision-making in the Nigerian banking sector. *International Journal of Scientific and Management Research*, 2022, 5, 12. DOI: 10.37502/IJSMR.2022.51215.
- [31] Mahalakshmi, V., Kulkarni, N., Kumar, K.P., Kumar, K.S., Sree, D.N. and Durga, S. The role of implementing Artificial Intelligence and Machine Learning technologies in the financial services industry for creating competitive intelligence. *Materials Today: Proceedings*, 2022, 56, 2252-2255. DOI: 10.1016/j.matpr.2021.11.577.
- [32] Pandey, T.N., Jagadev, A.K., Mohapatra, S.K. and Dehuri, S. Credit risk analysis using machine learning classifiers. 2017 International Conference on Energy, Communication, Data Analytics and Soft Computing (ICECDS). IEEE, 2017, 1850-1854. DOI: 10.1109/ICECDS.2017.8389769.
- [33] Wang, L., Cheng, Y., Gu, X. and Wu, Z. Design and optimization of big data and machine learning-based risk monitoring system in financial markets. *arXiv preprint arXiv:2407.19352*, 2024. DOI: 10.48550/arXiv.2407.19352.
- [34] Yang, M., Lim, M.K., Qu, Y., Ni, D. and Xiao, Z. Supply chain risk management with machine learning technology: A literature review and future research directions. *Computers & Industrial Engineering*, 2023, 175, 108859. DOI: 10.1016/j.cie.2022.108859.
- [35] Xiao, Y., Wu, J., Lin, Z. and Zhao, X. A deep learning-based multi-model ensemble method for cancer prediction. *Computer Methods and Programs in Biomedicine*, 2018, 153, 1-9. DOI: 10.1016/j.cmpb.2017.09.005.
- [36] Ahmed, K., Sachindra, D., Shahid, S., Iqbal, Z., Nawaz, N. and Khan, N. Multi-model ensemble predictions of precipitation and temperature using machine learning algorithms. *Atmospheric Research*, 2020, 236, 104806. DOI: 10.1016/j.atmosres.2019.104806.
- [37] Khamparia, A., Singh, A., Anand, D., Gupta, D., Khanna, A., Kumar, N.A. and Tan, J. A novel deep learning-based multi-model ensemble method for the prediction of neuromuscular disorders. *Neural Computing and Applications*, 2020, 32, 11083-11095. DOI: 10.1007/s00521-018-3896-0. ([EBSCO OpenURL][3])
- [38] Yu, J., Cai, Z., He, P., Xie, G. and Ling, Q. Multi-model ensemble learning method for human expression recognition. *arXiv preprint arXiv:2203.14466*, 2022. DOI: 10.48550/arXiv.2203.14466.
- [39] Zhang, Y., Fu, K., Sun, H., Sun, X., Zheng, X. and Wang, H. A multi-model ensemble method based on convolutional neural networks for aircraft detection in large remote sensing images. *Remote Sensing Letters*, 2018, 9(1), 11-20. DOI: 10.1080/2150704X.2017.1378452.
- [40] Li, X., Chen, H., Xu, L., Mo, Q., Du, X. and Tang, G. Multi-model fusion stacking ensemble learning method for

- the prediction of berberine by FT-NIR spectroscopy. *Infrared Physics & Technology*, 2024, 137, 105169. DOI: 10.1016/j.infrared.2024.105169.
- [41] Ashfaq, A., Imran, A., Ullah, I., Alzahrani, A., Alheeti, K.M.A. and Yasin, A. Multi-model ensemble based approach for heart disease diagnosis. 2022 International Conference on Recent Advances in Electrical Engineering & Computer Sciences (RAEE & CS). IEEE, 2022, 1-8. DOI: 10.1109/RAEECS56511.2022.9954490.
- [42] Wang, Z., Dong, J. and Zhang, J. Multi-model ensemble deep learning method to diagnose COVID-19 using chest computed tomography images. *Journal of Shanghai Jiaotong University (Science)*, 2022, 1-11. DOI: 10.1007/s12204-021-2392-3.
- [43] Moro, S., Rita, P. and Cortez, P. Bank Marketing. UCI Machine Learning Repository, 2014. DOI: 10.24432/C5K306.
- [44] Zabor, E.C., Reddy, C.A., Tendulkar, R.D. and Patil, S. Logistic regression in clinical studies. *International Journal of Radiation Oncology Biology Physics*, 2022, 112(2), 271-277. DOI: 10.1016/j.ijrobp.2021.08.007.
- [45] Naseer, A. and Jalal, A. Pixels to precision: Features fusion and random forests over label-based segmentation. IEEE, 2023, 1-6. DOI: 10.1109/IBCAST59916.2023.10712869.
- [46] Liu, W., Chen, Z. and Hu, Y. XGBoost algorithm-based prediction of safety assessment for pipelines. *International Journal of Pressure Vessels and Piping*, 2022, 197, 104655. DOI: 10.1016/j.ijpvp.2022.104655.
- [47] Tian, L., Feng, L., Yang, L. and Guo, Y. Stock price prediction based on LSTM and LightGBM hybrid model. *The Journal of Supercomputing*, 2022, 78(9), 11768-11793. DOI: 10.1007/s11227-022-04326-5.
- [48] Fan, Z., Gou, J. and Weng, S. A feature importance-based multi-layer CatBoost for student performance prediction. *IEEE Transactions on Knowledge and Data Engineering*, 2024. DOI: 10.1109/TKDE.2024.3393472.

Copyright© by the authors, Licensee Intelligence Technology International Press. The article is an open-access article distributed under the terms and conditions of the Creative Commons Attribution-ShareAlike 4.0 (CC BY-SA).